



جامعة الشهيد حمّة لخضر - الوادي
Université Echahid Hamma Lakhdar - El-Oued
كلية العلوم الدقيقة – قسم الفيزياء

Instrumentation (3^{ème} année physique)

MEHELLOU Saïd

mehellou-said@univ-eloued.dz

2023

Préface

Niveau : 3^{ème} année physique

Semestre : 06

Matière : Instrumentation

Objectifs :

L'étudiant va apprendre l'essentiel sur le fonctionnement des différents systèmes de détection de radiations (détecteurs de radiations) qui repose sur les propriétés de l'interaction des rayonnements incidents avec le matériau constituant le détecteur.

Connaissances préalables :

L'étudiant doit avoir des connaissances de base sur les propriétés des rayonnements, la physique des interactions rayonnement-matière et les circuits électriques de base. L'essentiel de certains de ces sujets sont passés en revue dans ce cours.

Contenu de la matière :

PROPRIETES GENERALES DES DETECTEURS OPTIQUES

1. Modèle simplifié d'un détecteur ;
2. Fonctionnement en mode courant et en mode pulse ;
 - 2.1 Cas où RC est petit ($\tau \ll t_c$);
 - 2.2 Cas où RC est grand ($\tau \gg t_c$);
3. Spectres des hauteurs de pulses ;
4. Courbes de comptage et plateaux ;
5. Résolution en énergie ;
6. Efficacité de détection ;
7. Temps mort ;
 - 7.1 Modèles pour comportement du temps mort ;
 - 7.2 Méthodes de mesure du temps mort.

Annexe I : Détecteur de radiations à gaz ;

Annexe II : Détecteur de radiations scintillant ;

Annexe III : Détecteur de radiations à solide (semi-conducteur).

Tale des matières

Avant-propos	06
Introduction générale.....	07
Chapitre 01 : Rayonnements et sources de rayonnements	
1. Introduction.....	09
2. Définition du rayonnement.....	09
3. Classification des rayonnements.....	09
3.1 Rayonnements non ionisants.....	09
3.2 Rayonnements ionisants.....	09
3.2.1 Classification des rayonnements ionisants selon leurs modes d'ionisation.....	09
3.2.1.1 Rayonnements directement ionisants.....	10
3.2.1.2 Rayonnements indirectement ionisants.....	10
3.2.2 Classification des rayonnements ionisants selon leurs natures.....	10
3.2.2.1 Rayonnements particuliers.....	11
3.2.2.2 Rayonnements électromagnétiques.....	11
4. Natures des rayonnements et les sources de rayonnements	12
4.1 Natures des rayonnements	12
4.2 Sources de rayonnements.....	13
4.2.1 Sources d'électron rapide.....	13
4.2.2 Sources de particules lourdes chargées.....	15
4.2.3 Sources de rayonnements électromagnétiques	15
4.2.4 Sources des neutrons.....	19
Chapitre 02: Interaction rayonnement-matière	
1. Introduction.....	22
2. Interaction des particules lourdes chargées.....	23
2.1. Nature de l'interaction.....	23
2.2 Puissance d'arrêt	24
2.3 Caractéristiques de perte d'énergie.....	25
2.3.1 Courbe de Bragg	25
2.4 Parcours des particules.....	25
2.4.1 Définition du parcours.....	25
2.4.2 Écarts des parcours.....	26

2.4.3 Temps d'arrêt.....	27
2.5 Perte d'énergie dans les absorbeurs minces.....	28
2.6 Comportement des fragments de fission.....	29
2.7 Émission d'électrons secondaires de surface.....	29
3. Interaction des électrons rapides.....	30
3.1 Perte d'énergie spécifique.....	30
3.2 Parcours de l'électron et courbes de transmission.....	31
3.2.1. Absorption des électrons mono-énergétiques.....	31
3.2.2 Absorption des particules Beta	32
3.2.3 Rétrodiffusion.....	33
3.3 Interactions des positrons.....	33
4. Interactions des rayons Gamma.....	34
4.1 Mécanismes d'interaction.....	34
4.1.1 Absorption photoélectrique	34
4.1.2 Diffusion Compton.....	35
4.1.3 Production de paires électron-positon.....	35
4.1.4 Diffusion cohérente.....	36
4.2 Atténuation des rayons gamma.....	37
4.2.1 Coefficients d'atténuation.....	37
4.2.2 Epaisseur de la masse absorbante.....	38
5. Interactions des neutrons.....	39
5.1 Propriétés générales.....	39
5.2 Interactions des neutrons lents.....	40
5.3 Interactions des neutrons rapides	40
5.4 Sections efficaces des neutrons.....	41
6. Exposition aux rayonnements et dose.....	42
6.1 Exposition aux rayons Gamma.....	42
6.2 Dose absorbée	43
6.3 Équivalent de dose.....	43
6.4 Conversion fluence-dose.....	44
6.5 Unités de dose en International Commission on Radiation Protection (ICRP).....	45
6.6 Quantités de dose opérationnelle.....	46

Chapitre 03 : Propriétés générales des détecteurs de rayonnements

1. Introduction.....	47
2. Modèle simplifié d'un détecteur.....	47
3. Modes de fonctionnement d'un détecteur.....	49
3.1 Mode courant.....	49
3.2 Mode de tension quadratique moyenne.....	52
3.3 Mode d'impulsion.....	53
3.3.1 Cas où RC est très faible ($\tau \ll t_c$).....	54
3.3.2 Cas où RC est très grande ($\tau \gg t_c$).....	54
4. Spectres des hauteurs des impulsions.....	57
5. Courbes de comptage et plateaux.....	59
6. Résolution en énergie.....	61
7. Efficacité de détection.....	66
8. Temps mort.....	70
8.1 Modèles de comportement du temps mort.....	71
8.2 Méthodes de mesure du temps mort.....	74
8.3 Statistiques des pertes du temps mort.....	78
8.4 Pertes de temps mort provenant de sources pulsées.....	79
Annexe I : Détecteur de radiations à gaz	98
Annexe II : Détecteur de radiations scintillant	142
Annexe III : Détecteur de radiations à solide (semi-conducteur)	179
Bibliographie.....	251
Abstract.....	252
المخلص.....	253

Avant-propos

Ce cours est destiné principalement aux étudiants du 1^{er} cycle universitaire (spécialité physique), il a pour objectif d'enseigner les éléments de base des détecteurs de radiations tout en se focalisant sur l'étude des propriétés générales de ces systèmes. Puisque, le principe de fonctionnement des détecteurs de radiations repose sur les propriétés de l'interaction des rayonnements incidents avec le matériau constituant le détecteur, il a fallu le débiter par un rappel des notions élémentaires de la physique des interactions rayonnement-matière. De même, le cours dans sa globalité est conçu de manière à développer une bonne compréhension des mécanismes physiques des détecteurs qui sont à la base de leurs applications dans le domaine de la détection de radiations.

Introduction générale

Les effets constatés lors des premières recherches sur les rayons X et les matières radioactives restent à la base de la détection telle qu'on la pratique de nos jours. Les instruments de détection ont été maintes fois perfectionnés au cours des années et de nombreuses recherches ont contribué à ces progrès.

C'est sur différents principes tels que la scintillation et l'ionisation de diverses substances qu'a reposé le fonctionnement de plusieurs types de détecteurs. La détection des rayonnements par la lumière scintillante induite dans une substance qui par la suite converti l'énergie de rayonnement incident en fluorescence instantanée, a été un des premiers procédés de détection.

Par la suite, des tubes électroniques photomultiplicateurs ont été associés aux scintillateurs pour convertir les faibles éclats lumineux en impulsions électriques qu'on peut compter électroniquement.

L'emploi du principe d'ionisation dans un appareil de détection exige que les charges produites par le rayonnement puissent se déplacer sous l'influence d'un champ électrique. Les gaz remplissent facilement cette condition et les détecteurs à ionisation de la première génération contenaient un gaz.

Bien que des détecteurs à gaz soient encore d'usage courant, les détecteurs à ionisation classiques avec remplissage de gaz ont été abandonnés pour presque tous les travaux de haute précision. On n'a pas tardé à se rendre compte qu'il y avait grand intérêt à remplacer le gaz par un milieu solide. Les solides ont en effet une densité environ 1000 fois supérieure à celle des gaz et permettent d'employer des détecteurs beaucoup moins encombrants. Ces détecteurs de rayonnements mesurent l'ionisation produite dans des matériaux semi-conducteurs et qui sont l'équivalent solide des chambres remplies de gaz. Les détecteurs semi-conducteurs ne tardèrent pas à se développer, aujourd'hui, c'est surtout le silicium qu'on emploie comme semi-conducteur dans les détecteurs à diode, alors que le germanium est plutôt utilisé pour les détecteurs à migration d'ions. L'emploi des détecteurs semi-conducteurs se heurte à une difficulté pratique, ils doivent être stockés et fonctionner à très basse température.

La maîtrise des détecteurs de rayonnements est inséparable de la physique des interactions rayonnement-matière. Le principe de fonctionnement des systèmes de détection repose dans la plupart des cas sur les propriétés de l'interaction des rayonnements détectés avec le matériau constituant le détecteur.

La compréhension du fonctionnement d'un détecteur nécessite donc de connaître ces processus d'interaction. C'est la raison pour laquelle ce cours regroupe à la fois les notions fondamentales

concernant les rayonnements et leurs sources ainsi que leurs interactions avec la matière en plus des différentes familles de détecteurs et leurs propriétés générales.

Les détecteurs introduits dans ce cours sont de structure simple. Ils constituent la base de systèmes de détection plus complexes qui sont notamment présentés dans d'autres cours (Master).

La compréhension des processus d'interaction rayonnement-matière nécessite des connaissances approfondies sur les rayonnements et les sources de rayonnements, ceci fera l'objet du premier chapitre.

L'étude du fonctionnement des différents types de détecteurs de rayonnements passe par une maîtrise préalable de la physique des interactions rayonnement-matière, le deuxième chapitre présentera les concepts de base de ces processus.

Pour mieux traiter les mécanismes physiques responsables du fonctionnement des détecteurs de rayonnements, le troisième chapitre abordera en détail les propriétés générales de ces systèmes avant de passer en revue les principales familles des détecteurs de radiations respectivement en annexes I, II, et III.

Chapitre : 01

Rayonnements et sources de rayonnements

1. Introduction

Le rayonnement est une forme d'énergie qui peut se déposer en totalité ou en partie dans un milieu approprié et produire ainsi un effet. La détection et la mesure du rayonnement reposent sur la détection et la mesure de ses effets dans un milieu, et l'histoire de l'apparition des détecteurs de rayonnement est étroitement liée à la découverte des rayonnements et de leurs effets.

2. Définition du rayonnement

On appelle rayonnement ou radiation l'émission ou la transmission d'énergie sous la forme d'ondes électromagnétiques ou de particules, et selon la valeur de l'énergie on peut les classer en deux catégories.

3. Classification des rayonnements

Selon leurs effets sur la matière (selon leurs énergies), les rayonnements peuvent être classés en deux catégories.

- Rayonnements non ionisants ;
- Rayonnements ionisants.

3.1 Rayonnements non ionisants

Ce sont des rayonnements qui renferment les ondes électromagnétiques les moins énergétiques, dont la longueur d'onde est supérieure à 100 nm.

3.2 Rayonnements ionisants

Les rayonnements ionisants consistent en particules, y compris des photons, qui arrachent des électrons à des atomes et à des molécules. Toutefois, certains rayonnements d'énergie relativement faible, comme les rayons ultraviolets, peuvent être ionisants dans des conditions particulières. Pour les distinguer des rayonnements qui provoquent toujours l'ionisation, on définit un seuil arbitraire d'énergie, en général 10 kilo-électronvolts (keV), à partir duquel les rayonnements sont dits ionisants. Exemples : les particules chargées (e^- , e^+ , protons, α).

3.2.1 Classification des rayonnements ionisants selon leurs modes d'ionisation

Les rayonnements ionisants se distinguent en deux catégories :

- Rayonnements directement ionisants ;
- Rayonnements indirectement ionisants.

3.2.1.1 Rayonnements directement ionisants

Ce sont toutes les particules chargées ; électrons, protons, particules α et ions lourds. Ils cèdent directement leurs énergies aux électrons atomiques du milieu (ou moins probablement au noyau) à travers l'interaction coulombienne.

3.2.1.2 Rayonnements indirectement ionisants

Ce sont les rayonnements neutres ; neutrons, photons γ et X. Ils transfèrent leurs énergies en deux étapes :

- En première étape, le rayonnement libère une particule chargée dans le milieu (les photons relâchent des électrons ou des positons. Les neutrons libèrent des protons ou ions lourds).
- Dans la deuxième étape, les particules chargées libérées déposent leurs énergies au milieu à travers l'interaction coulombienne.

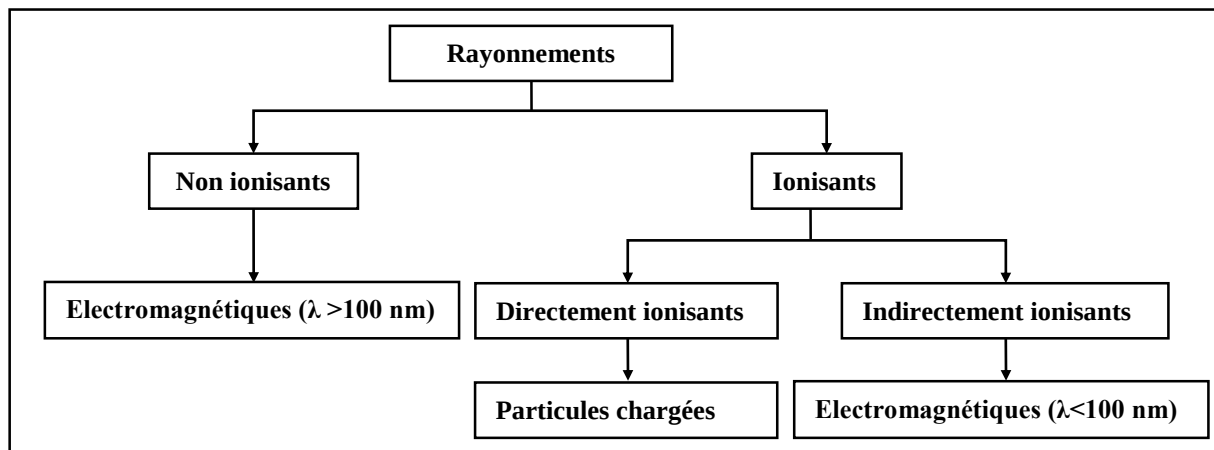


Fig. 01 : Classification des rayonnements selon leurs modes d'ionisation

3.2.2 Classification des rayonnements ionisants selon leurs natures

Les rayonnements ionisants peuvent être classifiés en deux groupes selon leur nature :

- Rayonnements particuliers ;
- Rayonnements électromagnétiques.

3.2.2.1 Rayonnements particuliers :

Ils représentent les particules possédant une masse au repos. Elles peuvent être :

- Chargées : électrons, positons, protons, ions (particules chargées lourdes). Elles interagissent avec le milieu par interactions coulombiennes, principalement avec les électrons atomiques (Interaction plus probables) et moins fréquemment avec les noyaux.
- Neutres : les neutrons ; lors de leur passage dans le milieu, ils perdent de l'énergie, soit, par ralentissement ou par collisions sur les noyaux du milieu traversé (interaction favorable pour les

neutrons rapides), soit par capture neutronique et réactions nucléaires (interactions très probables pour les neutrons lents et thermiques).

3.2.2.2 Rayonnements électromagnétiques :

Ce sont les photons γ résultant de transitions nucléaires, les rayons X caractéristiques émis lors des transitions des électrons en orbites atomiques et les rayonnements de freinage (Bremsstrahlung) envoyés lors des interactions coulombiennes des électrons rapides avec le noyau atomique (décélération des électrons rapides au voisinage du noyau). Les rayonnements électromagnétiques interagissent avec la matière par différents processus : effet photoélectrique, diffusion Compton, production des paires électron-positron, ...

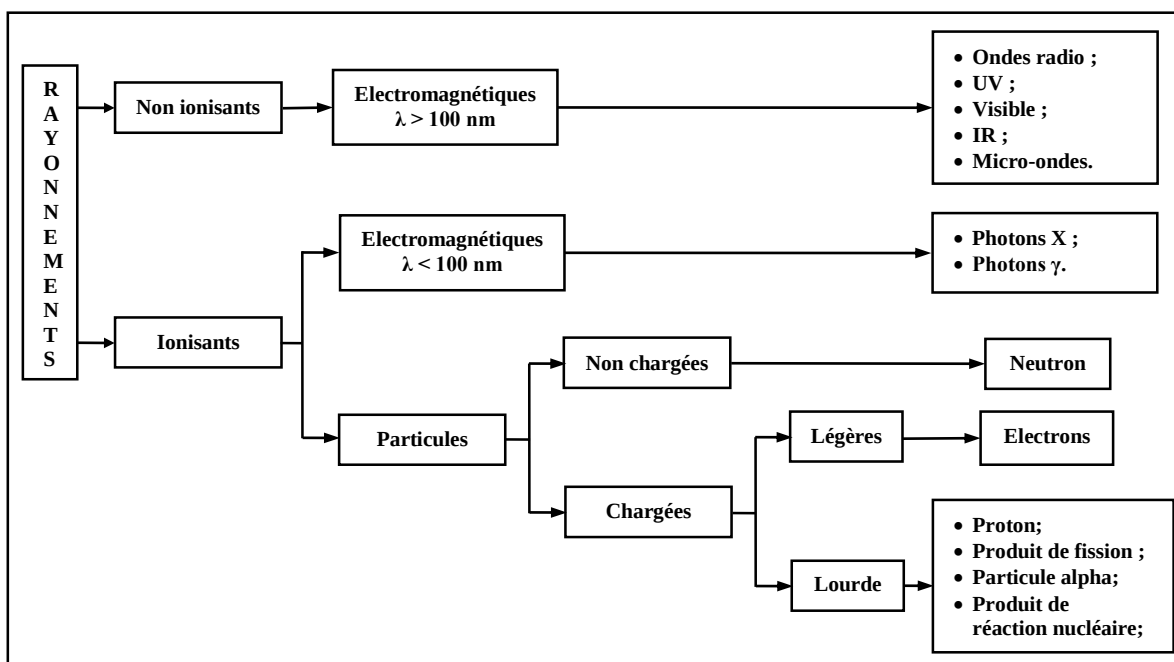


Fig. 02 : Classification des rayonnements selon leur nature

4. Natures des rayonnements et les sources de rayonnements

4. 1 Natures des rayonnements

Les rayonnements concernés dans cette étude proviennent de processus atomiques ou nucléaires. Ils sont commodément classés en quatre types généraux comme suit:

- Rayonnement particulaire chargé :
 - Électrons rapides ;
 - Particules chargées lourdes.
- Rayonnement non chargé :
 - Rayonnement électromagnétique ;
 - Neutrons.

Les électrons rapides comprennent les particules « bêta » (positives ou négatives) émises lors de la désintégration nucléaire, ainsi que les électrons énergétiques produits par tout autre processus. Les particules chargées lourdes désignent une catégorie qui englobe tous les ions énergétiques avec une masse d'une unité de masse atomique ou plus, tels que les particules alpha, les protons, les produits de fission ou les produits de nombreuses réactions nucléaires. Le rayonnement électromagnétique d'intérêt comprend les rayons X émis lors du réarrangement des couches d'électrons des atomes et les rayons gamma qui proviennent des transitions à l'intérieur du noyau lui-même. Les neutrons générés dans divers procédés nucléaires constituent une grande catégorie, qui est souvent subdivisée en sous-catégories de neutrons lents et de neutrons rapides.

La gamme d'énergie d'intérêt s'étend d'environ 10 eV à 20 MeV. (Les neutrons lents sont techniquement une exception mais sont inclus en raison de leur importance.) La limite d'énergie inférieure est définie par l'énergie minimale requise pour produire l'ionisation dans les matériaux typiques par le rayonnement. Les radiations avec une énergie supérieure à ce minimum sont classées comme rayonnements ionisants.

Les rayonnements diffèrent par leur "dureté" ou leur capacité de pénétrer à des épaisseurs dans le matériau. Cette propriété est très importante pour déterminer la forme physique des sources de rayonnement. Les radiations douces, telles que les particules alpha ou les rayons X de faible énergie, ne pénètrent qu'à de petites épaisseurs dans le matériau. Les particules bêta sont généralement plus pénétrantes, ainsi que les rayons gamma ou les neutrons qui sont des radiations plus dures.

Dans ce qui suit, on mettra principalement l'accent sur les sources de rayonnements, qui sont susceptibles d'être intéressantes soit pour l'étalonnage et les essais des détecteurs de rayonnement décrits dans les chapitres suivants ou comme objets des mesures elles-mêmes. Le rayonnement d'origine naturelle est une source supplémentaire importante.

4. 2 Sources de rayonnements

4. 2. 1 Sources d'électron rapide

A. Désintégration bêta

La source la plus courante d'électrons rapides dans les mesures de rayonnement est un radio-isotope qui se désintègre par émission bêta moins. Le processus est décrit schématiquement par :



Où X et Y sont les espèces nucléaires initiale et finale, et $\bar{\nu}$ est l'antineutrino. Les neutrinos et les antineutrinos ont une probabilité d'interaction extrêmement faible avec la matière, ils sont donc indétectables.

Le noyau de recul Y apparaît avec une très faible énergie de recul, qui est habituellement inférieure au seuil d'ionisation, et par conséquent, il ne peut pas être détecté par des moyens conventionnels. Ainsi, le seul rayonnement ionisant significatif produit par la désintégration bêta est l'électron rapide ou la particule bêta elle-même.

La plupart des désintégrations bêta peuplent un état excité du noyau produit, de sorte que les rayons gamma de désexcitation ultérieurs sont émis avec les particules bêta dans de nombreuses sources bêta courantes.

Chaque transition de désintégration bêta spécifique est caractérisée par une énergie de désintégration fixe ou valeur Q. Comme l'énergie du noyau de recul est pratiquement nulle, cette énergie est partagée entre la particule bêta et le neutrino "invisible". La particule bêta apparaît donc avec une énergie qui varie d'une désintégration à l'autre et peut aller de zéro à "l'énergie bêta du point final (end-point-energy)", qui est numériquement égale à la valeur Q. Un spectre d'énergie bêta représentatif est illustré à la figure (03). La valeur Q pour une désintégration donnée est normalement donnée en supposant que la transition a lieu entre les états fondamentaux des noyaux parent et fils.

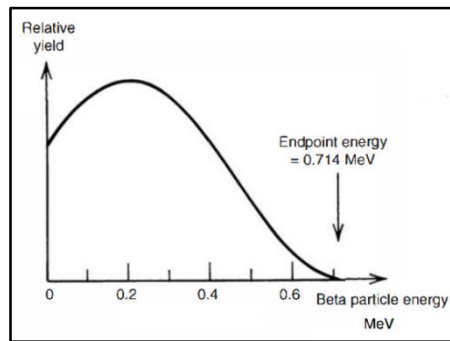


Fig. 03 : Distribution d'énergie des particules « bêta »

B. Conversion interne

Le processus de conversion interne commence par un état nucléaire excité, qui peut être formé par un processus antérieur, souvent les désintégrations bêta d'une espèce parent. La méthode courante de désexcitation consiste à émettre un photon gamma. Pour certains états excités, l'émission gamma peut être quelque peu inhibée et l'alternative de conversion interne peut devenir significative. L'énergie d'excitation nucléaire (E_{ex}) est transférée directement à l'un des électrons orbitaux de l'atome. Cet électron apparaît alors avec une énergie donnée par :

$$E_{e^-} = E_{ex} - E_b \quad (02)$$

Où (E_b) est l'énergie de liaison de l'électron dans la couche électronique d'origine.

C. Électrons Auger

Les électrons Auger sont à peu près l'analogue des électrons de conversion interne lorsque l'énergie d'excitation provient de l'atome plutôt que du noyau. Un processus antérieur (tel que la capture d'électron) peut laisser l'atome avec un vide dans une couche d'électrons normalement complète. Ce vide est souvent comblé par un électron provenant de l'une des couches externes de l'atome avec l'émission d'un photon X typique. Alternativement, l'énergie d'excitation de l'atome peut être transférée directement à l'un des électrons externes, provoquant son éjection de l'atome. Cet électron est appelé électron Auger et apparaît avec une énergie donnée par la différence entre l'énergie d'excitation atomique d'origine et l'énergie de liaison de la couche d'où l'électron a été éjecté.

4. 2. 2 Sources de particules lourdes chargées

A. Désintégration alpha

Les noyaux lourds sont énergétiquement instables face à l'émission spontanée d'une particule alpha (ou noyau ^4He). Le processus de désintégration est décrit schématiquement par :



Où X et Y sont les espèces nucléaires initiale et finale.

B. Fission spontanée

Le processus de fission est la seule source spontanée de particules énergétiques chargées lourdes de masse supérieure à celle de la particule alpha. Les fragments de fission sont donc largement utilisés dans l'étalonnage et le test des détecteurs destinés généralement à une application de mesures d'ions lourds.

Tous les noyaux lourds sont, en principe, instables vis-à-vis de la fission spontanée en deux fragments plus légers. Pour tous, excepté les noyaux extrêmement lourds, le processus est inhibé par la grande barrière de potentielle qui doit être surmontée lors de la distorsion du noyau de sa forme quasi sphérique d'origine. La fission spontanée n'est donc pas un processus significatif sauf pour certains isotopes transuraniens de très grand nombre de masse.

4. 2. 3 Sources de rayonnements électromagnétiques

A. Rayons gamma suite de la désintégration bêta

Le rayonnement gamma est émis par les noyaux excités lors de leur transition vers les niveaux nucléaires inférieurs. Dans la plupart des sources de laboratoire pratiques, les états nucléaires excités sont créés lors de la désintégration d'un parent radionucléide.

B. Rayonnement d'annihilation

Lorsque le noyau parent subit une désintégration « β^+ », un rayonnement électromagnétique supplémentaire est généré dont l'origine réside du sort des positrons émis lors du processus de désintégration primaire. Les positrons ne parcourent généralement que quelques millimètres avant de perdre leur énergie cinétique, l'encapsulation inhérente autour de la source est souvent suffisamment épaisse pour arrêter complètement les positrons. Lorsque leur énergie est très faible, ils se recombinaient avec des électrons dans les matériaux absorbants en cours d'annihilation. Le positon et l'électron d'origine disparaissent et sont remplacés par deux photons électromagnétiques (0,511 MeV) de directions opposées connus sous le nom de rayonnement d'annihilation. Ce rayonnement est ensuite superposé à tout rayonnement gamma pouvant être émis lors de la désintégration ultérieure du produit filiation.

C. Rayons gamma suite à des réactions nucléaires

Si des rayons gamma avec des énergies supérieures à celles disponibles à partir des isotopes bêta-actifs sont nécessaires, un autre processus doit conduire à la population d'états nucléaires plus élevés. Une possibilité est la réaction nucléaire :



Où le noyau produit ${}^{12}\text{C}$ est dans un état excité. Sa désintégration donne naissance à un photon gamma d'une énergie de 4,44 MeV.

D. Bremsstrahlung (rayonnement de freinage)

Lorsque des électrons rapides interagissent avec la matière, une partie de leur énergie est convertie en rayonnement électromagnétique sous forme de Bremsstrahlung. La fraction de l'énergie électronique convertie en Bremsstrahlung accroît avec l'augmentation de l'énergie des électrons et est maximale pour les matériaux absorbants à numéro atomique élevé. Le processus est important dans la production de rayons X à partir de tubes à rayons X conventionnels.

En plus du Bremsstrahlung, des rayons X typiques sont également produits lorsque des électrons rapides traversent un absorbeur. Par conséquent, les spectres des tubes à rayons X ou d'autres sources de Bremsstrahlung montrent également des raies d'émission de rayons X typiques superposées au spectre de Bremsstrahlung continu.

E. Rayons X typiques

Si les électrons orbitaux d'un atome sont perturbés par rapport à leur configuration normale par un processus d'excitation, l'atome peut exister dans un état excité pendant une courte période. Les

électrons ont une tendance naturelle à se réorganiser pour ramener l'atome à son niveau énergétique le plus bas ou à l'état fondamental dans un temps qui est typiquement d'une nanoseconde ou moins dans un matériau solide. L'énergie libérée lors de la transition de l'état excité à l'état fondamental prend la forme d'un photon X typique dont l'énergie est donnée par la différence d'énergie entre l'état initial et l'état final.

D'autres processus physiques différents peuvent conduire à la population d'états atomiques excités qui sont à l'origine des rayons X typiques. En général, l'énergie des photons X typiques est fixée par les énergies de liaison atomique.

➤ **Excitation par désintégration radioactive**

Dans le processus de désintégration nucléaire de la capture d'électrons, la charge nucléaire est diminuée d'une unité par la capture d'un électron orbital, le plus souvent un électron K. L'atome résultant a toujours le bon nombre d'électrons orbitaux, mais le processus de capture crée un vide dans l'une des couches internes. Lorsque ce vide est rempli, des rayons X sont générés et qui sont caractéristiques de l'élément produit. La désintégration peut peupler soit l'état fondamental, soit un état excité dans le noyau produit, de sorte que les rayons X typiques peuvent également être accompagnés de rayons gamma provenant de la désexcitation nucléaire ultérieure.

La conversion interne est un autre processus nucléaire qui peut conduire à des rayons X typiques. Comme défini précédemment dans ce chapitre, la conversion interne se traduit par l'éjection d'un électron orbital de l'atome en laissant derrière lui un vide. Encore une fois, ce sont les électrons K qui sont le plus facilement convertis et, par conséquent, les rayons X typiques de la série K sont les plus importants. La désexcitation gamma de l'état nucléaire est toujours un processus concurrent à la conversion interne, les sources de radio-isotopes de ce type émettent généralement des rayons gamma en plus des rayons X typiques. Les électrons de conversion peuvent également conduire à un continuum mesurable de Bremsstrahlung, en particulier lorsque leur énergie est élevée.

L'auto-absorption est un problème technique important dans la préparation des sources de rayons X radio-isotopes. Au fur et à mesure que l'épaisseur du dépôt de radio-isotopes augmente, le flux de rayons X émergent de sa surface se rapproche d'une valeur limite car seuls les atomes proches de la surface peuvent contribuer aux photons qui s'échappent.

➤ **Excitation par rayonnement externe**

Une source externe de rayonnement (rayons X, électrons, particules alpha, etc.) est amenée à heurter une cible, créant des atomes excités ou ionisés dans la cible. Étant donné que de nombreux atomes ou ions excités dans la cible se dés excitent ensuite à l'état fondamental par l'émission de rayons X typiques, la cible peut servir de source de ces rayons X. L'énergie des rayons X émis dépend

du choix du matériau cible. Les cibles à faible numéro atomique produisent des rayons X typiques doux, et les cibles à Z élevé produisent des rayons plus durs ou des rayons X plus énergétiques. Le rayonnement incident doit avoir une énergie supérieure à l'énergie maximale du photon émis par la cible, car les états excités conduisant au rayonnement atomique correspondant à la transition doit être peuplée par le rayonnement incident.

A titre d'exemple, le rayonnement incident peut consister en des rayons X générés dans un tube à rayons X classique. Ces rayons X peuvent alors interagir dans la cible par absorption photoélectrique, et la désexcitation ultérieure des ions cibles crée leur spectre de rayons X typique. Dans ce cas, le processus est appelé fluorescence X. Bien que le spectre caractéristique des rayons X peut être perturbé par des photons diffusés par le faisceau de rayons X incident.

Une autre méthode d'excitation de la cible consiste à utiliser un faisceau d'électrons externe. Pour les cibles de faible numéro atomique, des potentiels d'accélération de quelques milliers de volts sont nécessaires de sorte que des sources d'électrons relativement compactes peuvent être utilisées. Dans le cas d'une excitation électronique, le spectre des rayons X typiques de la cible sera perturbé par le spectre des rayonnements de freinage continu également généré par les interactions des électrons incidents avec la cible. Pour les cibles minces, les photons de Bremsstrahlung sont préférentiellement émis vers la direction avant, tandis que les rayons X typiques sont émis de façon isotrope.

Le rayonnement incident peut également être constitué de particules chargées lourdes. Encore une fois, les interactions de ces particules dans la cible donneront naissance aux atomes excités nécessaires à l'émission de rayons X typiques.

F. Rayonnement synchrotron

Une autre forme de rayonnement est produite lorsqu'un faisceau d'électrons énergétiques est dévié sur une orbite circulaire. D'après la théorie électromagnétique, une petite fraction de l'énergie du faisceau est rayonnée au cours de chaque cycle du faisceau. Considéré à l'origine comme une nuisance par les concepteurs d'accélérateurs à haute énergie, ce rayonnement synchrotron s'est révélé ultérieurement être une forme très utile de rayonnement électromagnétique. Lorsqu'il est extrait de l'accélérateur dans une direction tangentielle à l'orbite du faisceau, le rayonnement apparaît comme un faisceau intense et hautement directionnel de photons dont l'énergie peut s'étendre de la lumière visible (quelques eV) aux énergies des rayons X (10^4 eV). Les monochromateurs peuvent être utilisés pour produire des faisceaux de photons presque mono-énergétiques avec une intensité très élevée.

4.2.4 Sources des neutrons

Bien que les noyaux créés avec une énergie d'excitation supérieure à l'énergie de liaison des neutrons puissent se désintégrer par émission de neutrons, ces états hautement excités ne sont pas produits à la suite d'un quelconque processus pratique de désintégration radioactive. Par conséquent, les sources isotopiques pratiques de neutrons n'existent pas dans le même sens que les sources de rayons gamma sont disponibles à partir de nombreux noyaux différents peuplés par la désintégration bêta. Les choix possibles pour les sources de neutrons radio-isotopiques sont beaucoup plus limités et reposent soit sur la fission spontanée, soit sur des réactions nucléaires pour lesquelles la particule incidente est le produit d'un processus de désintégration conventionnel.

A. Fission spontanée

De nombreux nucléides lourds transuraniens ont une probabilité appréciable de désintégration par fission spontanée. Plusieurs neutrons rapides sont rapidement émis lors de chaque événement de fission, de sorte qu'un échantillon d'un tel radionucléide peut être une source de neutrons isotopiques simple et pratique. D'autres produits du processus de fission sont les produits de fission lourds décrits précédemment, les rayons gamma de fission rapide, et l'activité bêta et gamma des produits de fission accumulés dans l'échantillon. Lorsqu'il est utilisé comme source de neutrons, l'isotope est généralement encapsulé dans un récipient suffisamment épais pour que seuls les neutrons rapides et les rayons gamma émergent de la source.

B. Sources de radio-isotopes (alpha, n)

Étant donné que les particules alpha énergétiques sont disponibles à partir de la désintégration directe d'un certain nombre de radionucléides pratiques, il est possible de fabriquer une petite source de neutrons autonome en mélangeant un isotope émetteur alpha avec un matériau cible approprié. Plusieurs matériaux cibles différents peuvent conduire à (alpha, n) réactions pour les énergies des particules alpha qui sont facilement disponibles dans la désintégration radioactive. Le rendement maximal en neutrons est obtenu lorsque le béryllium est choisi comme cible et que des neutrons sont produits par la réaction :



D'autres réactions induites par les particules alpha ont été parfois utilisées comme sources de neutrons, mais toutes ont un rendement de neutrons sensiblement inférieur par unité d'activité alpha par rapport à la réaction au béryllium.

C. Sources de photo-neutrons

Certains émetteurs de rayons gamma radio-isotopes peuvent également être utilisés pour produire des neutrons lorsqu'ils sont combinés avec un matériau cible approprié. Les sources de photo-neutrons résultantes sont basées sur la fourniture d'une énergie d'excitation suffisante à un noyau cible par absorption d'un photon de rayons gamma pour permettre l'émission d'un neutron libre. Seuls deux noyaux cibles, ^9Be et ^2H , sont de toute importance pratique pour les sources de photo-neutrons radio-isotopiques. Les réactions correspondantes peuvent s'écrire :

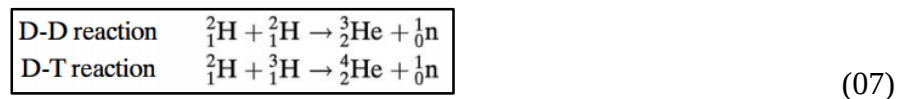


Un photon gamma avec une énergie au-dessus du seuil est nécessaire pour rendre les réactions énergétiquement possibles, de sorte que seuls les rayons gamma d'énergie relativement élevée peuvent être appliqués. Un avantage des sources de photo-neutrons est que si les rayons gamma sont mono-énergétiques, les neutrons sont également presque mono-énergétiques.

Le principal inconvénient des sources de photo-neutrons est qu'il faut utiliser de très grandes activités de rayons gamma pour produire des sources de neutrons d'intensité attractive.

D. Réactions de particules chargées accélérées

Étant donné que les particules alpha sont les seules particules chargées lourdes avec un faible Z facilement disponibles à partir de radio-isotopes, les réactions impliquant des protons incidents, des deutérons, etc. doivent s'appuyer sur des particules artificiellement accélérées. Deux des réactions les plus courantes de ce type pour produire des neutrons sont :



Du fait que la barrière coulombienne entre le deutéron incident et le noyau cible lumineux est relativement petite, les deutérons n'ont pas besoin d'être accélérés à une énergie très élevée afin de créer un rendement de neutrons significatif. Ces réactions sont largement exploitées dans des "générateurs de neutrons" dans lesquels les ions deutérium sont accélérés par un potentiel d'environ 100-300 kV.

D'autres réactions induites par des particules chargées qui impliquent par exemple une cible avec un numéro atomique élevé sont également utilisées pour la génération de neutrons. Pour ces réactions, des grands accélérateurs sont nécessaires pour produire des particules chargées avec des énergies élevées nécessaires à la production de neutrons.

Chapitre 02

Interaction rayonnement-matière

1. Introduction

Le fonctionnement de tout détecteur de rayonnement dépend essentiellement de la manière dont le rayonnement à détecter interagit avec le matériau du détecteur lui-même. La compréhension de la réponse d'un type particulier de détecteur doit donc reposer sur la connaissance des mécanismes fondamentaux par lesquels les rayonnements interagissent et perdent leur énergie dans la matière.

Avant d'entamer le sujet de ce chapitre qui est l'interaction rayonnement-matière, il convient de classer les quatre grandes catégories de rayonnements introduites au chapitre (01) dans le tableau suivant :

Tab. 01 : Catégories des rayonnements

Radiations particulières chargées		Radiations non chargées
Particules chargées lourdes (Distance caractéristique : 10^{-5} m)	←	Neutrons (Portée caractéristique : 10^{-1} m)
Électrons rapides (Distance caractéristique : 10^{-3} m)	←	Rayons X et rayons Gamma (Portée caractéristique : 10^{-1} m)

Les radiations particulières chargées qui, en raison de la charge électrique portée par la particule, interagissent par la force coulombienne avec les électrons présents dans tout milieu qu'ils traversent. Les radiations non chargées ne sont donc pas soumises à la force coulombienne, ces rayonnements doivent d'abord subir une interaction qui altère radicalement les propriétés du rayonnement incident en une seule interaction. L'interaction entraîne le transfert total ou partiel de l'énergie du rayonnement incident aux électrons ou aux noyaux des atomes constitutifs, ou aux particules chargées produites par les réactions nucléaires. Si l'interaction ne se produit pas dans le détecteur, ces rayonnements non chargés (par exemple, les neutrons ou les rayons gamma) peuvent traverser complètement le volume du détecteur sans révéler le moindre indice qu'ils étaient là.

Les flèches dans le tableau (01) indiquent les résultats de telles interactions. Un rayon X ou gamma, par les processus décrits dans ce chapitre, peut transférer toute ou une partie de son énergie aux électrons du milieu. Les électrons secondaires résultants présentent une similitude étroite avec les rayonnements d'électrons rapides (tels que les particules « bêta ») vues au chapitre (01). Les dispositifs conçus pour détecter les rayons gamma sont conçus pour favoriser de telles interactions et pour arrêter complètement les électrons secondaires résultants afin que toute leur énergie puisse contribuer au signal de sortie du détecteur. En revanche, les neutrons peuvent interagir

de manière à produire des particules chargées lourdes secondaires, qui servent alors de base au signal du détecteur.

Le tableau (01) contient également l'ordre de grandeur de la distance caractéristique de pénétration ou la longueur de trajet moyenne (portée ou libre parcours moyen) dans les solides pour les rayonnements d'énergie typiques dans chaque catégorie.

2. Interaction des particules lourdes chargées

2.1. Nature de l'interaction

Les particules chargées lourdes, telles que la particule alpha, interagissent avec la matière principalement par les forces coulombiennes entre leur charge positive et la charge négative des électrons orbitaux dans les atomes absorbants. Bien que des interactions de la particule avec les noyaux (comme dans la diffusion de Rutherford ou les réactions induites par les particules alpha) soient également possibles, de telles interactions se produisent que rarement et ils ne sont normalement pas significatifs dans la réponse des détecteurs de rayonnement. Au lieu de cela, les détecteurs de particules chargées doivent s'appuyer sur les résultats des interactions avec les électrons pour leur réponse.

En pénétrant dans n'importe quel milieu absorbant, la particule chargée interagit immédiatement et simultanément avec de nombreux électrons. Dans une telle interaction, l'électron ressent une impulsion de la force coulombienne attractive lorsque la particule passe à son voisinage. Selon la proximité de la rencontre, cette impulsion peut être suffisante soit pour élever l'électron vers une couche plus élevée dans l'atome absorbant (excitation) ou pour arracher complètement l'électron de l'atome (ionisation). L'énergie qui est transférée à l'électron doit se faire aux dépens de la particule chargée, et sa vitesse est donc diminuée à la suite de l'interaction. L'énergie maximale qui peut être transférée d'une particule chargée de masse m avec une énergie cinétique E à un électron de masse m_0 en une seule collision est de $4Em_0/m$, soit environ 1/500 de l'énergie de la particule par nucléon. Comme il s'agit d'une petite fraction de l'énergie totale, la particule primaire doit perdre son énergie dans de nombreuses interactions de ce type lors de son passage à travers un absorbeur. À tout moment, la particule interagit avec de nombreux électrons, de sorte que l'effet net est de diminuer continuellement sa vitesse jusqu'à ce que la particule soit arrêtée.

Les chemins représentatifs empruntés par les particules chargées lourdes dans leur processus de ralentissement sont représentés schématiquement dans la figure (01). Sauf à fin leur trajectoire, les pistes ont tendance à être assez droite parce que la particule n'est pas fortement déviée par une seule rencontre, et les interactions se produisent dans toutes les directions simultanément. Les particules chargées sont donc caractérisées par une plage définie dans un matériau absorbant donné. La plage,

à définir plus précisément ci-dessous, représente une distance au-delà de laquelle aucune particule ne pénétrera.

Les produits des interactions des particules chargées dans l'absorbeur sont soit des atomes excités, soit des paires d'ions. Chaque paire d'ions est composée d'un électron libre et de l'ion positif correspondant d'un atome absorbant dont un électron a été totalement arraché. Les paires d'ions ont une tendance naturelle à se recombiner pour former des atomes neutres, mais dans certains types de détecteurs, cette recombinaison est éliminée de sorte que les paires d'ions peuvent être utilisées comme base de la réponse du détecteur.

Lors des interactions particulièrement rapprochées, un électron peut subir une impulsion suffisamment importante pour qu'après avoir quitté son atome parent, il puisse encore avoir une énergie cinétique suffisante pour créer d'autres ions.

Ces électrons énergétiques sont parfois appelés rayons delta et représentent un moyen indirect par lequel l'énergie des particules chargées est transférée au milieu absorbant. Sous conditions typiques, la majorité de la perte d'énergie de la particule chargée se reproduit via ces rayons delta.

La portée des rayons delta est toujours petite par rapport à la portée de la particule énergétique incidente, de sorte que l'ionisation se forme toujours à proximité de la piste primaire. À l'échelle microscopique, un effet de ce processus est que les paires d'ions n'apparaissent normalement pas comme des ionisations simples espacées de manière aléatoire, mais il y a une tendance à former de nombreux "amas" de multiples paires d'ions répartis le long de la trajectoire de la particule.

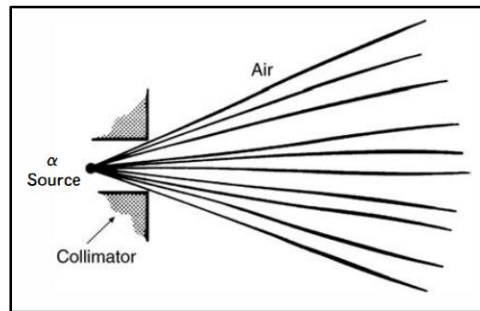


Fig. 01 Chemins représentatifs empruntés par les particules chargées lourdes dans leur processus de ralentissement

2.2 Puissance d'arrêt

Le pouvoir d'arrêt linéaire S pour les particules chargées dans un absorbeur donné est simplement défini comme le différentiel des pertes d'énergie pour cette particule dans le matériau divisée par la valeur correspondante du différentiel de la longueur du chemin :

$$S = -\frac{dE}{dx} \quad (01)$$

La valeur de $-dE/dx$ le long d'une trajectoire de particules est également appelée sa perte d'énergie spécifique ou, plus simplement, son "taux" de perte d'énergie.

Pour les particules avec un état de charge donné, S augmente à mesure que la vitesse des particules diminue.

2.3 Caractéristiques de perte d'énergie

2.3.1 Courbe de Bragg

La courbe représentant les pertes d'énergie spécifiques le long de la trajectoire d'une particule chargée est connue sous le nom de courbe de Bragg, figure (02).

La courbe associée montrant les variations de $-dE/dx$ en fonction de l'énergie des particules chargées lourdes illustrent l'énergie à laquelle la prise de charge par l'ion devient significative. Les particules chargées avec le plus grand nombre de charges nucléaires commencent à capter des électrons au début de leur processus de ralentissement. Notez que dans un absorbeur en aluminium, les ions « hydrogène » (protons) à charge unique montrent de forts effets de captage de charge en dessous d'environ 100 keV, mais les ions ^3He à double charge montrent des effets équivalents à environ 400 keV.

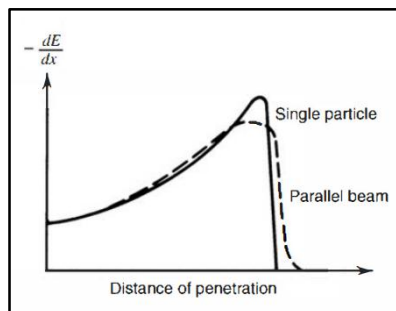


Fig. 02 Perte d'énergie spécifique le long d'une trajectoire alpha

2.4 Parcours des particules

2.4.1 Définition du parcours

Afin de quantifier la définition du parcours des particules, nous nous référons à l'expérience illustrée dans la figure (03). Une source collimatée de particules alpha mono-énergétiques comptées par un détecteur après passage dans un absorbeur d'épaisseur variable. Pour les particules alpha, les résultats sont également représentés sur la figure (03). Pour de petites valeurs de l'épaisseur de l'absorbeur, le seul effet est de provoquer une perte d'énergie des particules alpha dans l'absorbeur lors de leur passage.

Puisque les trajectoires à travers l'absorbeur sont assez droites, le nombre total qui atteint le détecteur reste le même. Aucune réduction du nombre de particules alpha n'a lieu jusqu'à ce que

l'épaisseur de l'absorbeur approche la longueur de la trajectoire la plus courte dans le matériau absorbant. L'augmentation de l'épaisseur s'arrête ainsi que le nombre de particules alpha, et l'intensité du faisceau détecté chute rapidement à zéro.

Le parcours des particules alpha dans le matériau absorbant peut-être déterminée à partir de cette courbe de plusieurs manières. Le parcours moyen est défini comme l'épaisseur de l'absorbeur qui réduit le nombre de particules alpha à la moitié de sa valeur en l'absence de l'absorbeur. Cette définition est la plus couramment utilisée dans les tableaux des valeurs numériques du parcours. Une autre version qui apparaît dans la littérature est le parcours extrapolé, qui est obtenu en extrapolant la partie linéaire de la fin de la courbe de transmission jusqu'à zéro.

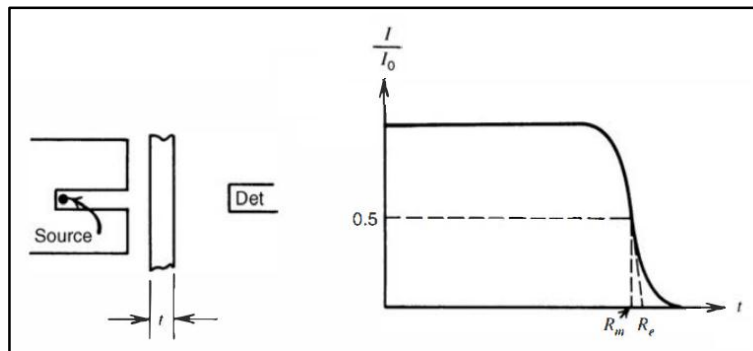


Fig. 03 Expérience de transmission de particules alpha ;

- (I) est le nombre de particules alpha détectées à travers une épaisseur d'absorbeur t ;
- (I_0) est le nombre particules alpha détectées sans absorbeur ;
- (R_m) et (R_e) sont respectivement le parcours moyen et le parcours extrapolé

2.4.2 Écarts des parcours

Les particules chargées sont également sujettes à un écart de parcours, défini comme la fluctuation de la longueur du trajet pour des particules individuelles de même énergie initiale. Les mêmes facteurs stochastiques qui conduisent à la dispersion de l'énergie à une distance de pénétration donnée entraînent également un chemin total longueurs légèrement différent pour chaque particule. Pour les particules chargées lourdes telles que les protons ou les particules alphas, la dispersion s'élève à quelques pour cent du parcours moyen. Le degré de dispersion est mis en évidence par la nette chute à la fin de la courbe de transmission moyenne tracée à la figure (02). La différenciation de cette courbe conduit à un pic dont la largeur souvent prise comme mesure quantitative de l'importance du décalage du parcours pour les particules et l'absorbant utilisés dans la mesure.

2.4.3 Temps d'arrêt

Le temps nécessaire pour arrêter une particule chargée dans un absorbeur peut être déduit de sa portée et de sa vitesse moyenne. Pour les particules non relativistes de masse m et d'énergie cinétique E , la vitesse est :

$$v = \sqrt{\frac{2E}{m}} \quad (02)$$

Si nous supposons que la vitesse moyenne des particules lors de leur ralentissement est $\langle v \rangle = K v$, où v est évalué à l'énergie initiale, alors le temps d'arrêt T peut être calculé à partir du parcours R comme :

$$T = \frac{R}{\langle v \rangle} = \frac{R}{Kc} \sqrt{\frac{mc^2}{2E}} \quad (03)$$

Si la particule était uniformément décélérée, alors $\langle v \rangle$ serait donné par $v/2$ et K serait de 0,5.

Cependant, les particules chargées perdent généralement de l'énergie à un rythme plus rapide vers la fin de leur plage et K devrait prendre une valeur un peu plus élevée. En supposant $K = 0,60$, le temps d'arrêt peut être estimé comme suit :

$$T \cong 1.2 \times 10^{-7} R \sqrt{\frac{m_A}{E}} \quad (04)$$

Où T est en secondes, R en mètres, m_A en unités de masse atomique et E en MeV. Cette approximation devrait être raisonnablement précise pour les particules chargées lourdes (protons, particules alpha, etc.) sur une grande partie de la gamme d'énergie. Elle ne doit pas cependant être utilisée pour les particules relativistes telles que les électrons rapides.

En utilisant des valeurs de plage typiques, les temps d'arrêt calculés à partir de cette équation pour les particules chargées sont de quelques picosecondes dans les solides ou les liquides et de quelques nanosecondes dans les gaz. Ces temps sont généralement suffisamment petits pour être négligés pour tous les détecteurs de rayonnement à réponse très rapide.

2.5 Perte d'énergie dans les absorbeurs minces

Pour les absorbeurs minces (ou détecteurs) qui sont pénétrés par une particule chargée donnée, l'énergie déposée dans l'absorbeur peut être calculée à partir de :

$$\Delta E = - \left(\frac{dE}{dx} \right)_{\text{avg}} t \quad (05)$$

Où t est l'épaisseur de l'absorbeur et $(-dE/dx)_{\text{avg}}$ est la puissance d'arrêt linéaire moyennée sur l'énergie de la particule dans l'absorbeur. Si la perte d'énergie est faible, la puissance d'arrêt ne varie pas beaucoup et peut être approximée par sa valeur à l'énergie des particules incidentes. Les valeurs tabulaires de dE/dx pour un certain nombre de particules chargées différentes dans une variété de milieux absorbants sont données dans plusieurs références.

Pour les épaisseurs d'absorbeur à travers lesquelles la perte d'énergie n'est pas faible, il n'est pas simple d'obtenir une valeur moyenne correctement pondérée $(-dE/dx)_{avg}$ directement à partir de ces données. Dans ces cas, il est plus facile d'obtenir l'énergie déposée d'une manière qui utilise des données parcourt-énergie. La base de cette méthode est la suivante : Soit R_1 le parcourt complet de la particule incidente d'énergie E_0 dans le matériau absorbant. En soustrayant l'épaisseur de l'absorbeur t de R_1 , on obtient une valeur R_2 qui représente le parcourt de ces particules alpha qui émergent de la surface opposée de l'absorbeur. En trouvant l'énergie correspondant à R_2 , on obtient l'énergie des particules chargées transmises E_t . L'énergie déposée ΔE est alors donnée simplement par $E_0 - E_t$. Ces étapes sont illustrées dans la figure (04) :

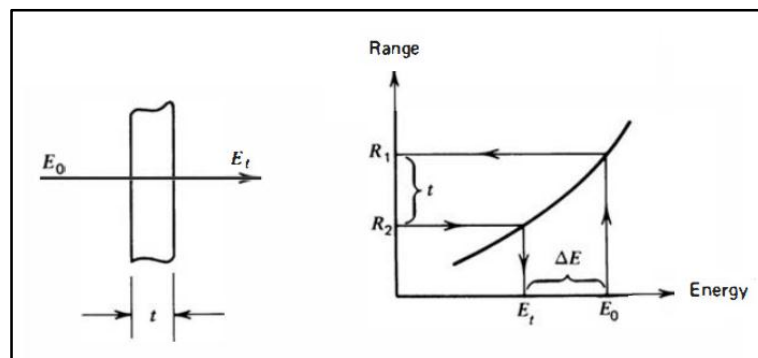


Fig. 04 Méthode de détermination de l'énergie des particules, déposée dans un absorbeur (cas de grandes pertes)

La procédure est basée sur l'hypothèse que les trajectoires des particules chargées sont parfaitement linéaires dans l'absorbeur, et la méthode ne s'applique pas dans les situations où la particule peut être déviée de manière significative (comme pour les électrons rapides).

2.6 Comportement des fragments de fission

Les fragments lourds produits à la suite de neutrons induits ou par la fission spontanée de noyaux lourds sont des particules énergétiques chargées avec des propriétés quelque peu différentes de celles discutées jusqu'à ce point. Parce que les fragments commencent par être dépouillés de nombreux électrons, leur très grande charge effective entraîne une perte d'énergie spécifique supérieure à celle rencontrée avec tout autre rayonnement discuté dans ce texte. Cependant, comme l'énergie initiale est également très élevée, la portée d'un fragment de fission typique est environ la moitié de celle d'une particule alpha de 5 MeV.

Une caractéristique importante d'un parcourt de fragment de fission est le fait que la perte d'énergie spécifique $-dE/dx$ diminue à mesure que la particule perd de l'énergie dans l'absorbeur. Ce comportement est en contraste marqué avec les particules plus légères, telles que les particules alpha

ou les protons, et est le résultat de la diminution continue de la charge effective portée par le fragment lorsque sa vitesse est réduite.

Le captage des électrons commence immédiatement au début du parcours, et donc le facteur z au numérateur de l'équation suivante chute continuellement,

$$\boxed{-\frac{dE}{dx} = \frac{4\pi e^4 z^2}{m_0 v^2} NB} \quad (06)$$

La diminution résultante de $-dE/dx$ est suffisamment importante pour surmonter l'augmentation qui accompagne normalement une réduction de vitesse. Pour des particules avec un état de charge initial beaucoup plus faible, comme la particule alpha, le captage d'électrons ne devient significatif que vers la fin du parcours.

2.7 Émission d'électrons secondaires de surface

Comme les particules chargées perdent leur énergie cinétique pendant le processus de ralentissement, de nombreux électrons de l'absorbeur reçoivent une fraction d'énergie suffisante pour parcourir une courte distance à partir de la trajectoire d'origine de la particule. Ceux-ci comprennent les rayons delta mentionnés précédemment, dont l'énergie est suffisamment élevée pour ioniser d'autres atomes absorbants, ainsi que des électrons dont l'énergie cinétique de quelques eV seulement inférieure à l'énergie minimale pour une ionisation ultérieure. Si une particule chargée énergétique est incidente sur une surface solide ou en émerge, certains de ces électrons peuvent migrer vers la surface et avoir suffisamment d'énergie pour s'échapper. Le terme électron secondaire est classiquement appliqué à ces électrons de faible énergie qui s'échappent. Remarque : le même terme est souvent utilisé pour caractériser les électrons beaucoup plus énergétiques formés dans les interactions de rayons gamma, type très différent d'électron secondaire.

3. Interaction des électrons rapides

Par rapport aux particules chargées lourdes, les électrons rapides perdent leur énergie à un rythme inférieur et suivent un chemin beaucoup plus tortueux à travers les matériaux absorbants. Une série de pistes provenant d'une source d'électrons mono-énergétiques pourrait apparaître comme dans la figure (05) :

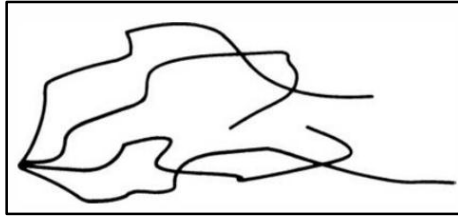


Fig. 05 Trajectoires des électrons rapides dans un matériau absorbant

De grandes déviations dans la trajectoire des électrons sont désormais possibles car leur masse est égale à celle des électrons orbitaux avec lesquels ils interagissent, et une fraction beaucoup plus importante de leur énergie peut être perdue en une seule rencontre. De plus, des interactions électron-nucléaire, qui peuvent brusquement changer la direction des électrons, se produisent parfois.

3.1 Perte d'énergie spécifique

L'expression qui décrit la perte d'énergie spécifique due à l'ionisation et à l'excitation "pertes collisionnelles" pour les électrons rapides est notée par :

$$\boxed{-\left(\frac{dE}{dx}\right)_c} \quad (07)$$

Les électrons diffèrent également des particules chargées lourdes par le fait que l'énergie peut être perdue par des processus radiatifs ainsi que par des interactions coulombiennes. Ces pertes radiatives prennent la forme de Bremsstrahlung ou de rayonnement électromagnétique, qui peuvent émaner de n'importe quelle position le long de la trajectoire de l'électron. D'après la théorie classique, toute charge doit émettre de l'énergie lorsqu'elle est accélérée, et les déviations de l'électron dans ses interactions avec l'absorbeur correspondent à une telle accélération.

La perte d'énergie spécifique linéaire par ce processus radiatif est notée par :

$$\boxed{-\left(\frac{dE}{dx}\right)_r} \quad (08)$$

Pour les types de particules et les gammes d'énergie considérées, seuls les électrons rapides peuvent avoir un rendement significatif de Bremsstrahlung. Le rendement des particules chargées lourdes est négligeable, tandis que les pertes radiatives sont plus importantes pour les électrons à hautes énergies et pour les matériaux absorbants de grand numéro atomique. Pour les énergies typiques des électrons, l'énergie photonique de freinage moyenne est assez faible et est donc normalement réabsorbée assez près de son point d'origine. Dans certains cas, cependant, la fuite de Bremsstrahlung peut influencer la réponse de petits détecteurs.

Le pouvoir d'arrêt linéaire total des électrons est la somme des pertes collisionnelles et radiatives :

$$\frac{dE}{dx} = \left(\frac{dE}{dx}\right)_c + \left(\frac{dE}{dx}\right)_r \quad (09)$$

Le rapport des pertes d'énergie spécifique est donné approximativement par :

$$\frac{(dE/dx)_r}{(dE/dx)_c} \cong \frac{EZ}{700} \quad (10)$$

Où E est l'énergie cinétique des particules et Z est le numéro atomique des atomes absorbants.

Pour les électrons (tels que les particules bêta ou électrons secondaires des interactions gamma), les énergies typiques sont inférieures à quelques MeV. Par conséquent, les pertes radiatives représentent toujours une petite fraction des pertes d'énergie dues à l'ionisation et à l'excitation et ne sont significatives que dans les matériaux absorbants à numéro atomique élevé.

3.2 Parcours de l'électron et courbes de transmission

3.2.1. Absorption des électrons mono-énergétiques

Une expérience d'atténuation du type discuté précédemment pour les particules alpha est esquissée sur la figure (06) pour une source d'électrons rapides mono-énergétiques. Même de faibles épaisseurs de l'absorbeur entraînent la perte de certains électrons du faisceau détecté car la diffusion de l'électron l'élimine efficacement du flux frappant le détecteur. Un tracé du nombre d'électrons détectés en fonction de l'épaisseur de l'absorbeur commence donc à chuter immédiatement et se rapproche progressivement de zéro pour les grandes épaisseurs de l'absorbeur. Les électrons qui pénètrent à la plus grande épaisseur de l'absorbeur seront ceux dont la direction initiale est la moins changée dans leur chemin à travers l'absorbeur.

Le concept de portée (parcourt) est moins précis pour les électrons rapides que pour les particules chargées lourdes, car la longueur totale du trajet des électrons est considérablement supérieure à la distance de pénétration suivant le vecteur vitesse, initial. Normalement, le parcours des électrons est tirée de la courbe de transmission, tel que celui donné à la figure 06, par extrapolation de la partie linéaire de la courbe vers zéro et représenter l'épaisseur de l'absorbeur requise pour s'assurer que presque aucun électron ne peut pénétrer toute l'épaisseur.

Pour une énergie équivalente, la perte d'énergie spécifique des électrons est bien inférieure à celle des particules chargées lourdes, de sorte que leur parcours dans les absorbeurs typiques est de centaines de fois supérieure. En tant qu'estimation, les parcours des électrons ont tendance à être d'environ 2 mm par Me V dans les matériaux à faible densité, ou d'environ 1 mm par MeV dans les matériaux de densité modérée.

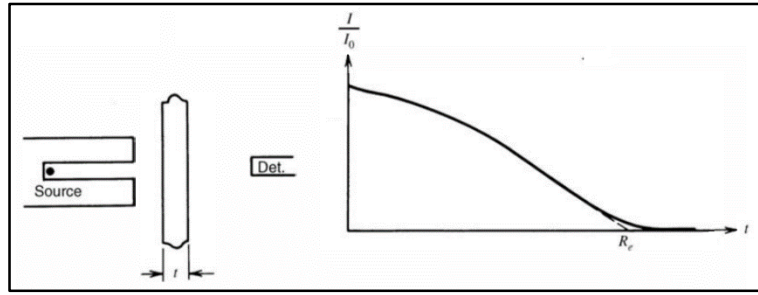


Fig. 06 Courbe de transmission des électrons mono-énergétiques
(R_e est le parcours extrapolé)

3.2.2 Absorption des particules Beta

La courbe de transmission des particules « bêta » émises par une source radio-isotopes, en raison de la distribution continue de leur énergie, diffère sensiblement de celle esquissée sur la figure 06 pour les électrons mono-énergétiques. Les particules « bêta » de faible énergie sont rapidement absorbées même dans de faibles épaisseurs d'absorbeur, de sorte que la pente initiale sur la courbe d'atténuation est beaucoup plus importante. Pour la majorité des spectres « bêta », la courbe a une forme quasi exponentielle et est donc presque linéaire. Ce comportement exponentiel n'est qu'une approximation empirique et n'a pas de base fondamentale comme l'atténuation exponentielle des rayons gamma. Un coefficient d'absorption n est parfois défini par :

$$\boxed{\frac{I}{I_0} = e^{-nt}} \quad (11)$$

Où :

I_0 : taux de comptage sans absorbeur ;

I : taux de comptage avec absorbeur ;

t : épaisseur de l'absorbeur.

3.2.3 Rétrodiffusion

Le fait que les électrons subissent souvent des déviations de large angle le long de leurs trajectoires conduit au phénomène de rétrodiffusion. Un électron incident à une surface d'un absorbeur peut subir une déviation suffisante pour qu'il réémerge de la surface par laquelle il est entré. Ces électrons rétrodiffusés ne déposent pas toute leur énergie dans le milieu absorbant et peuvent donc avoir un effet significatif sur la réponse des détecteurs destinés à mesurer l'énergie des électrons incidents. Les électrons qui rétrodiffusent dans la "fenêtre d'entrée" du détecteur ou la couche morte échapperont entièrement à la détection.

La rétrodiffusion est plus importante pour les électrons de faible énergie incidente et les absorbeurs de numéro atomique élevé. La rétrodiffusion peut également influencer le rendement

apparent des sources radio-isotopes de particules bêta ou d'électrons de conversion. Si la source est déposée sur un support épais, les électrons initialement émis dans ce support peuvent rétrodiffuser et réémerger de la surface de la source.

3.3 Interactions des positrons

Les forces coulombiennes qui constituent le principal mécanisme de perte d'énergie pour les électrons et les particules chargées lourdes sont présentes pour une charge positive ou négative sur la particule. L'interaction implique une force répulsive ou attractive entre la particule incidente et l'électron orbital, l'impulsion et le transfert d'énergie pour des particules de masse égale sont à peu près les mêmes. Par conséquent, les trajectoires des positrons dans un absorbeur sont similaires à celles des électrons, et leurs pertes d'énergie spécifique et leurs parcours sont à peu près les mêmes pour des énergies initiales égales.

Cependant, les positons diffèrent considérablement en ce sens qu'un rayonnement d'annihilation est généré à la fin du parcours des positrons. Comme ces photons de 0,511 Me V sont très pénétrants par rapport au parcours du positon, ils peuvent conduire au dépôt d'énergie loin de la trajectoire d'origine du positon.

4. Interactions des rayons Gamma

Bien qu'un grand nombre de mécanismes d'interaction possibles soient connus pour les rayons gamma avec la matière, seuls trois grands types jouent un rôle important dans les mesures de rayonnement : absorption photoélectrique, diffusion Compton et génération de paires. Tous ces processus conduisent au transfert partiel ou total de l'énergie des photons gamma en énergie d'électrons. Ils entraînent des changements brusques dans l'historique des photons gamma, le photon disparaît complètement ou est diffusé selon un large angle. Ce comportement est en remarquable contraste avec les particules chargées, qui ralentissent graduellement à travers des interactions continues et simultanées avec de nombreux atomes absorbants.

4.1 Mécanismes d'interaction

4.1.1 Absorption photoélectrique

Dans le processus d'absorption photoélectrique, un photon subit une interaction avec un atome absorbant dans laquelle le photon disparaît complètement. A sa place, un photoélectron énergétique est éjecté par l'atome de l'une de ses couches liées. L'interaction se fait avec l'atome dans son ensemble et ne peut pas avoir lieu avec des électrons libres. Pour les rayons gamma d'énergie suffisante, l'origine la plus probable du photoélectron est la couche la plus étroitement liée ou la couche K de l'atome. Le photoélectron apparaît avec une énergie donnée par :

$$E_{e^-} = h\nu - E_b$$

(12)

Où E_b représente l'énergie de liaison du photoélectron dans sa couche d'origine. Pour des énergies de rayons gamma supérieures à quelques centaines de keV, le photoélectron emporte la majorité de l'énergie originale des photons.

En plus du photoélectron, l'interaction crée également un atome absorbant ionisé avec un vide dans l'une de ses couches liées. Ce vide est rapidement comblé par la capture d'un électron libre du milieu et/ou le réarrangement d'électrons d'autres couches de l'atome. Par conséquent, un ou plusieurs photons de rayons X caractéristiques peuvent également être générés. Bien que dans la plupart des cas ces rayons X soient réabsorbés à proximité du site d'origine par absorption photoélectrique impliquant des couches moins étroitement liées, leur migration et leur fuite possible des détecteurs de rayonnement peuvent influencer leur réponse. Dans de rares cas, l'émission d'un électron Auger peut remplacer le rayon X caractéristique en emportant l'énergie d'excitation atomique.

4.1.2 Diffusion Compton

Le processus d'interaction de la diffusion Compton a lieu entre le photon gamma incident et un électron dans le matériau absorbant. C'est le plus souvent le mécanisme d'interaction prédominant pour les énergies de rayons gamma typiques des sources de radio-isotopes.

Dans la diffusion Compton, le photon gamma entrant est dévié d'un angle θ par rapport à sa direction d'origine. Le photon transfère une partie de son énergie à l'électron (supposé initialement au repos), qui est alors appelé électron de recul. Puisque tous les angles de diffusion sont possibles, l'énergie transférée à l'électron peut varier de zéro à un grand taux de l'énergie des rayons gamma.

L'expression qui relie le transfert d'énergie et l'angle de diffusion pour une interaction donnée peut être simplement dérivée en écrivant des équations simultanées pour la conservation de l'énergie et de la quantité de mouvement. En utilisant les symboles définis dans la figure (07) :

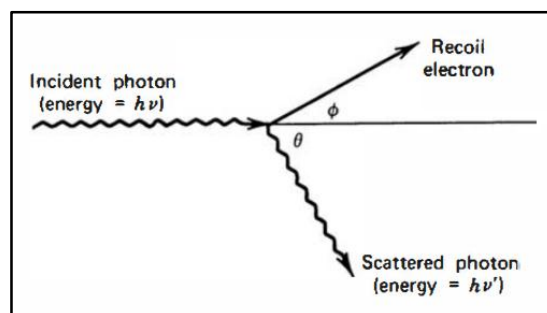


Fig. 07 Principe de la diffusion Compton

On peut montrer que :

$$\boxed{h\nu' = \frac{h\nu}{1 + \frac{h\nu}{m_0c^2}(1 - \cos \theta)}} \quad (13)$$

Où m_0c^2 est l'énergie de masse au repos de l'électron (0,511 MeV). Pour les faibles angles de diffusion θ , une faible quantité d'énergie est transférée. Une partie de l'énergie d'origine est toujours retenue par le photon incident, même à l'état extrême de $\theta = \pi$.

La probabilité de la diffusion Compton par atome de l'absorbeur dépend du nombre d'électrons disponibles comme cibles de diffusion et augmente donc linéairement avec Z.

4.1.3 Production de paires électron-positon

Si l'énergie des rayons gamma dépasse deux fois l'énergie de masse au repos d'un électron (1,02 MeV), le processus de production de paires est énergétiquement possible. Aux énergies des rayons gamma qui ne sont que quelques centaines de keV au-dessus de ce seuil, la probabilité de production de paires est faible. Cependant, ce mécanisme d'interaction devient prédominant lorsque l'énergie augmente dans la gamme de plusieurs MeV. Dans l'interaction (qui doit avoir lieu dans le champ coulombien d'un noyau), le photon gamma disparaît et est remplacé par une paire électron-positon. Toute l'énergie excédentaire transportée par le photon au-dessus des 1,02 MeV nécessaires à la création de la paire passe en énergie cinétique partagée par le positron et l'électron. Étant donné que le positron s'annihilera ensuite après avoir ralenti dans le milieu absorbant, deux photons d'annihilation sont normalement produits en tant que produits secondaires de l'interaction. Le sort ultérieur de ce rayonnement d'annihilation a un effet important sur la réponse des détecteurs de rayons gamma.

L'importance relative des trois processus décrits ci-dessus pour différents matériaux absorbants et énergies de rayons gamma est commodément illustrée à la figure (08). La ligne de gauche représente l'énergie à laquelle l'absorption photoélectrique et la diffusion Compton sont également probables en fonction du numéro atomique de l'absorbeur. La ligne de droite représente l'énergie à laquelle la diffusion Compton et la production de paires sont également probables. Trois zones sont ainsi définies sur le tracé au sein desquelles l'absorption photoélectrique, la diffusion Compton et la production de paires prédominent chacune.

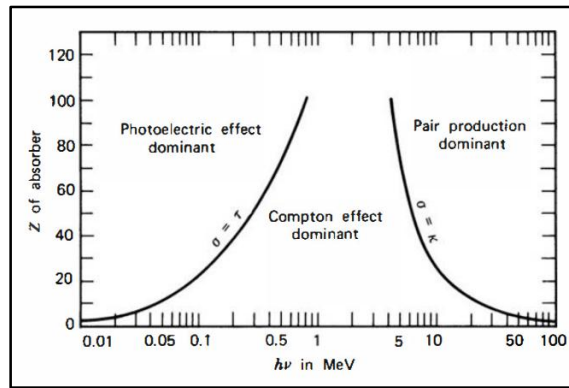


Fig. 08 Importance relative des trois grands types d'interaction gamma

4.1.4 Diffusion cohérente

En plus de la diffusion Compton, un autre type de diffusion peut se produire dans lequel le photon gamma interagit de manière cohérente avec tous les électrons d'un atome absorbant. Ce processus de diffusion cohérente ou de diffusion Rayleigh n'excite ni n'ionise l'atome, et le photon gamma conserve son énergie d'origine après l'événement de diffusion. Puisque pratiquement aucune énergie n'est transférée, ce processus est souvent négligé dans les discussions de base sur les interactions des rayons gamma. Cependant, la direction du photon est modifiée en diffusion cohérente, et des modèles complets de transport de rayons gamma doivent en tenir compte.

La probabilité de diffusion cohérente n'est significative que pour les faibles énergies de photons (généralement inférieures à quelques centaines de keV pour les matériaux courants) et est plus importante dans les absorbeurs à Z élevé.

L'angle de déviation moyen diminue avec l'augmentation de l'énergie, ce qui limite encore l'importance pratique de la diffusion cohérente aux basses énergies.

4.2 Atténuation des rayons gamma

4.2.1 Coefficients d'atténuation

Si nous imaginons à nouveau une expérience de transmission comme sur la figure (09), où les rayons gamma mono-énergétiques sont collimatés en un faisceau étroit et autorisés à frapper un détecteur après avoir traversé un absorbeur d'épaisseur variable, le résultat devrait être une simple atténuation exponentielle des rayons gamma comme également illustré à la figure (09).

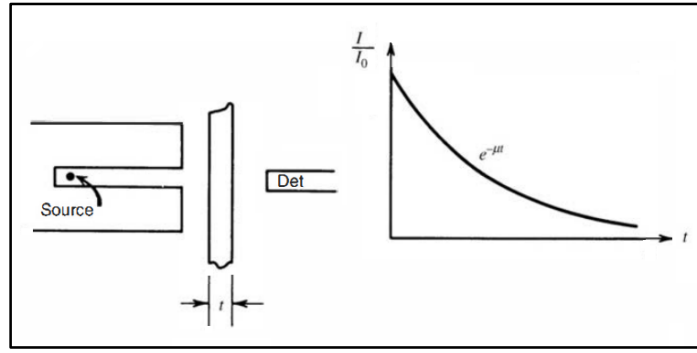


Fig. 09 Courbe de transmission (exponentielle) des rayons gamma

Chacun des processus d'interaction élimine le photon gamma du faisceau soit par absorption, soit par diffusion à partir de la direction du détecteur et peut être caractérisé par une probabilité fixe d'occurrence par unité de longueur de trajet dans l'absorbeur. La somme de ces probabilités est simplement la probabilité par unité de longueur de trajet que le photon gamma soit retiré du faisceau :

$$\mu = \tau(\text{photoelectric}) + \sigma(\text{Compton}) + \kappa(\text{pair}) \quad (14)$$

Et est appelé coefficient d'atténuation linéaire. Le nombre de photons transmis I est alors donné en termes de nombre sans absorbeur I_0 comme :

$$\frac{I}{I_0} = e^{-\mu t} \quad (15)$$

Les photons gamma peuvent également être caractérisés par leur libre parcours moyen λ défini comme la distance moyenne parcourue dans l'absorbeur avant qu'une interaction n'ait lieu. Sa valeur peut être obtenue à partir de :

$$\lambda = \frac{\int_0^{\infty} x e^{-\mu x} dx}{\int_0^{\infty} e^{-\mu x} dx} = \frac{1}{\mu} \quad (16)$$

Et est simplement l'inverse du coefficient d'atténuation linéaire. Les valeurs typiques de λ vont de quelques mm à des dizaines de cm dans les solides pour les énergies courantes des rayons gamma.

L'utilisation du coefficient d'atténuation linéaire est limitée par le fait qu'il varie avec la densité de l'absorbeur, même si le matériau de l'absorbeur est le même. Par conséquent, le coefficient d'atténuation de masse est beaucoup plus largement utilisé et est défini comme :

$$\boxed{\text{mass attenuation coefficient} = \frac{\mu}{\rho}} \quad (17)$$

Où ρ représente la densité du milieu. Pour une énergie gamma donnée, le coefficient d'atténuation massique ne change pas avec l'état physique d'un absorbeur donné. Par exemple, il en est de même pour l'eau qu'elle se présente sous forme liquide ou vapeur. Le coefficient d'atténuation massique d'un composé ou d'un mélange d'éléments peut être calculé à partir de :

$$\boxed{\left(\frac{\mu}{\rho}\right)_c = \sum_i w_i \left(\frac{\mu}{\rho}\right)_i} \quad (18)$$

Où les facteurs w_i représentent la fraction pondérale de l'élément i dans le composé ou le mélange.

4.2.2 Epaisseur de la masse absorbante

En termes de coefficient d'atténuation massique, la loi d'atténuation pour les rayons gamma prend désormais la forme :

$$\boxed{\frac{I}{I_0} = e^{-(\mu/\rho)\rho t}} \quad (19)$$

Le produit ρt , appelé épaisseur de masse de l'absorbeur, est désormais le paramètre significatif qui détermine son degré d'atténuation. Les unités d'épaisseur de masse ont toujours été mg/cm^2 . L'épaisseur des absorbeurs utilisés dans les mesures de rayonnement est donc souvent mesurée en épaisseur de masse plutôt qu'en épaisseur physique, car il s'agit d'une grandeur physique plus fondamentale.

L'épaisseur de masse est également un concept utile pour discuter la perte d'énergie des particules chargées et des électrons rapides. Pour des matériaux absorbants avec des rapports neutron/proton similaires, une particule rencontrera environ le même nombre d'électrons traversant des absorbeurs d'épaisseur de masse égale. Par conséquent, la puissance d'arrêt et la portée, lorsqu'elles sont exprimées en unités de ρt , sont à peu près les mêmes pour les matériaux qui ne diffèrent pas beaucoup en Z .

5. Interactions des neutrons

5.1 Propriétés générales

Comme les rayons gamma, les neutrons ne portent aucune charge et ne peuvent donc pas interagir dans la matière au moyen de la force coulombienne, qui domine les mécanismes de perte d'énergie pour les particules chargées et les électrons. Les neutrons peuvent également traverser

plusieurs centimètres de matière sans aucun type d'interaction et peuvent donc être totalement invisibles pour un détecteur. Lorsqu'un neutron subit une interaction, c'est avec un noyau du matériau absorbant. Du fait de l'interaction, le neutron peut soit totalement disparaître et être remplacé par un ou plusieurs rayonnements secondaires, soit l'énergie ou la direction du neutron est modifiée de manière significative.

Contrairement aux rayons gamma, les rayonnements secondaires résultant des interactions neutroniques sont presque toujours des particules chargées lourdes. Ces particules peuvent être produites soit à la suite de réactions nucléaires induites par des neutrons, soit elles peuvent être les noyaux du matériau absorbant lui-même, qui ont gagné de l'énergie à la suite de collisions de neutrons. La plupart des détecteurs de neutrons utilisent un certain type de conversion du neutron incident en particules chargées secondaires, qui peuvent ensuite être détecté directement.

Les probabilités relatives des différents types d'interactions neutroniques changent considérablement avec l'énergie neutronique. Dans une simplification un peu exagérée, nous diviserons les neutrons en deux catégories sur la base de leur énergie, soit les « neutrons rapides », soit les « neutrons lents », et discuterons séparément leurs propriétés d'interaction. La ligne de division sera à environ 0,5 eV, soit environ l'énergie de la chute brutale de la section efficace d'absorption dans le cadmium (l'énergie de coupure du cadmium).

5.2 Interactions des neutrons lents

Pour les neutrons lents, les interactions significatives comprennent la diffusion élastique avec les noyaux absorbants et un grand nombre de réactions nucléaires induites par les neutrons. En raison de la faible énergie cinétique des neutrons lents, très peu d'énergie peut être transférée au noyau lors de la diffusion élastique. Par conséquent, ce n'est pas une interaction sur laquelle peuvent se baser des détecteurs de neutrons lents. Cependant, les collisions élastiques ont tendance à être très probables et servent souvent à amener le neutron lent en équilibre thermique avec le milieu absorbant avant qu'un autre type d'interaction ne se produise. Une grande partie de la population dans le domaine d'énergie des neutrons lents se trouvera donc parmi ces neutrons thermiques qui, à température ambiante, ont une énergie moyenne d'environ 0,025 eV.

Les interactions neutroniques lentes d'importance réelle sont des réactions induites par les neutrons qui peuvent créer des rayonnements secondaires d'une énergie suffisante pour être détectés directement. Étant donné que l'énergie des neutrons entrants est si faible, toutes ces réactions doivent avoir une valeur Q positive pour être énergétiquement possibles. Dans la plupart des matériaux, la réaction de capture radiative [ou réaction (n, γ)] est la plus probable et joue un rôle important dans l'atténuation ou le blindage des neutrons. Les réactions de capture radiative peuvent être utiles dans la détection indirecte de neutrons à l'aide de feuilles d'activation, mais elles sont peu appliquées dans

les détecteurs actifs de neutrons car le rayonnement secondaire prend la forme de rayons gamma, également difficiles à détecter. Au lieu de cela, des réactions telles que (n, α) , (n, p) et $(n, \text{fission})$ sont beaucoup plus attrayantes car les rayonnements secondaires sont des particules chargées.

5.3 Interactions des neutrons rapides

La probabilité de la plupart des réactions induites par les neutrons potentiellement utiles dans les détecteurs diminue rapidement avec l'augmentation de l'énergie des neutrons. L'importance de la diffusion devient cependant plus grande, car le neutron peut transférer une quantité appréciable d'énergie en une seule collision.

Les rayonnements secondaires dans ce cas sont des noyaux de recul, qui ont capté une quantité détectable d'énergie provenant de collisions de neutrons. À chaque site de diffusion, le neutron perd de l'énergie et est ainsi modéré ou ralenti à une énergie inférieure. Le modérateur le plus efficace est l'hydrogène car le neutron peut perdre jusqu'à toute son énergie en une seule collision avec un noyau d'hydrogène. Pour les noyaux plus lourds, seul un transfert d'énergie partiel est possible.

Si l'énergie du neutron rapide est suffisamment élevée, une diffusion inélastique avec des noyaux peut avoir lieu dans laquelle le noyau de recul est élevé à l'un de ses états excités pendant la collision. Le noyau se désexcite rapidement, émettant un rayon gamma, et le neutron perd une plus grande fraction de son énergie qu'il ne le ferait dans une collision élastique équivalente. La diffusion inélastique et les rayons gamma secondaires qui en résultent jouent un rôle important dans le blindage des neutrons de haute énergie, mais constituent une complication indésirable dans la réponse de la plupart des détecteurs de neutrons rapides basés sur la diffusion élastique.

5.4 Sections efficaces des neutrons

Pour les neutrons d'énergie fixe, la probabilité par unité de longueur de parcours est une constante pour chacun des mécanismes d'interaction. Il est classique d'exprimer cette probabilité en termes de σ de section efficace par noyau pour chaque type d'interaction. La section transversale a des unités de surface et a traditionnellement été mesurée en unités de la grange (10^{-28} m^2). Par exemple, chaque espèce nucléaire aura une section efficace de diffusion élastique, une section efficace de capture radiative, etc., chacune de qui sera fonction de l'énergie des neutrons.

Lorsqu'elle est multipliée par le nombre de noyaux N par unité de volume, la section efficace σ est convertie en section efficace macroscopique Σ

$$\boxed{\Sigma = N\sigma} \quad (20)$$

Σ a maintenant les dimensions de l'inverse de la longueur. La section efficace macroscopique Σ a l'interprétation physique de la probabilité par unité de longueur de trajet pour un processus spécifique décrit par la section transversale "microscopique" σ .

Lorsque tous les processus sont combinés, en additionnant les sections transversales pour chaque interaction individuelle :

$$\Sigma_{\text{tot}} = \Sigma_{\text{scatter}} + \Sigma_{\text{rad. capture}} + \dots \quad (21)$$

Le Σ_{tot} résultant est la probabilité par unité de longueur de trajet que n'importe quel type d'interaction qui peut se produire.

Cette quantité a la même signification pour les neutrons que le coefficient d'absorption linéaire pour les rayons gamma défini précédemment. Si une expérience d'atténuation de faisceau étroit est effectuée pour les neutrons, comme esquissé précédemment pour les rayons gamma, le résultat sera le même : le nombre des neutrons détectés chutera de façon exponentielle avec l'épaisseur de l'absorbeur. Dans ce cas la relation d'atténuation s'écrit :

$$\frac{I}{I_0} = e^{-\Sigma_{\text{tot}} t} \quad (22)$$

Le libre parcours moyen des neutrons λ est, par analogie avec le cas des rayons gamma, donné par $1/\Sigma_{\text{tot}}$. Dans les matériaux solides, λ pour les neutrons lents peut être de l'ordre du centimètre ou moins, alors que pour les neutrons rapides, il est normalement des dizaines de centimètres.

6. Exposition aux rayonnements et dose

En raison de leur importance dans la protection du personnel dans les installations produisant des rayonnements et dans les applications médicales des rayonnements, les concepts d'exposition aux rayonnements et de dose jouent un rôle de premier plan dans les mesures de rayonnement. Dans les sections suivantes, nous introduisons les concepts fondamentaux qui sous-tendent les quantités et les unités d'importance dans ce domaine. Les définitions précises des grandeurs liées à la radioprotection évoluent continuellement et le lecteur est renvoyé aux publications de la Commission internationale des unités et mesures radiologiques (ICRU), de la Commission internationale de protection radiologique (ICRP) et du National Council on Radiation Protection des États-Unis (NCRP).

6.1 Exposition aux rayons Gamma

Le concept d'exposition aux rayons gamma a été introduit au début de l'histoire de la recherche sur les radionucléides. Défini uniquement pour les sources de rayons X ou gamma, un taux

d'exposition fixe existe en tout point de l'espace entourant une source d'intensité fixe. L'unité de base de l'exposition aux rayons gamma est définie en termes de charge dQ due à l'ionisation créée par les électrons secondaires (électrons négatifs et positrons) formés dans un volume élémentaire d'air et de masse dm , lorsque ces électrons secondaires sont complètement arrêtés dans l'air. La valeur d'exposition X est alors donnée par dQ/dm . L'unité d'exposition (en SI) aux rayons gamma est donc le coulomb par kilogramme (C/kg), qui n'a pas de nom particulier. L'unité historique a été le roentgen (R), défini comme l'exposition qui entraîne la génération d'une unité de charge électrostatique (environ $2,08 \times 10^9$ paires d'ions) par 0,001293 g (1 cm^3 à STP) d'air. Les deux unités sont liées par :

$$\boxed{1 \text{ R} = 2.58 \times 10^{-4} \text{ C/kg}} \quad (23)$$

L'exposition est donc définie en termes d'effet d'un flux donné de rayons gamma sur un volume d'essai d'air et n'est fonction que de l'intensité de la source intégrée sur le temps de mesure, de la géométrie entre la source et le volume d'essai, et toute atténuation des rayons gamma qui peuvent avoir lieu entre les deux. Sa mesure nécessite fondamentalement la détermination de la charge due à l'ionisation produite dans l'air dans des conditions spécifiques.

6.2 Dose absorbée

Deux matériaux différents, s'ils sont soumis à la même exposition aux rayons gamma, absorberont en général des quantités d'énergie différentes. Étant donné que de nombreux phénomènes importants, y compris des changements dans les propriétés physiques ou les réactions chimiques induites, on s'attendrait à ce que l'échelle de l'énergie absorbée par unité de masse du matériau, une unité qui mesure cette quantité est d'un intérêt fondamental. L'énergie moyenne absorbée par tout type de rayonnement par unité de masse de l'absorbeur est définie comme la dose absorbée. L'unité historique de dose absorbée a été le rad, défini comme 100 ergs/gramme.

Comme pour les autres unités de rayonnement historiques, le rad a été remplacé par son équivalent SI, le gray (Gy) défini comme 1 joule/kilogramme. Les deux unités sont donc simplement liées par :

$$\boxed{1 \text{ Gy} = 100 \text{ rad}} \quad (24)$$

La dose absorbée est une mesure raisonnable des effets chimiques ou physiques créés par une exposition donnée à un rayonnement dans un matériau absorbant. Des mesures minutieuses ont montré que la dose absorbée dans l'air correspondant à une exposition aux rayons gamma de 1 coulomb/kilogramme s'élève à 33,8 joules/kilogramme, soit 33,8 Gy.

6.3 Équivalent de dose

La gravité et la permanence des dommages biologiques créés par les rayonnements ionisants sont directement liés au taux local de dépôt d'énergie le long de la trajectoire des particules, connu sous le nom de transfert d'énergie linéaire L . Les rayonnements avec de grandes valeurs de L (comme les particules lourdes chargées) ont tendance à entraîner des dommages biologiques plus importants que ceux avec un L inférieur (comme les électrons), même si l'énergie totale déposée par unité de masse est la même.

Le concept d'équivalent de dose a donc été introduit pour quantifier plus adéquatement l'effet biologique probable d'une exposition donnée aux rayonnements. Une unité d'équivalent de dose est définie comme la quantité de tout type de rayonnement qui, lorsqu'elle est absorbée dans un système biologique, entraîne le même effet biologique qu'une unité de dose absorbée délivrée sous forme de rayonnement à faible L . L'équivalent de dose H pour un type de rayonnement donné est le produit de la dose absorbée D en un point du tissu et du facteur de qualité Q en ce point :

$$\boxed{H = DQ} \quad (25)$$

Le facteur de qualité augmente avec le transfert d'énergie linéaire L , pour les rayonnements d'électrons rapides, L est suffisamment faible pour que Q soit essentiellement égal à l'unité dans toutes les applications. Par conséquent, l'équivalent de dose est numériquement égal à la dose absorbée pour les particules « bêta » ou d'autres électrons rapides. Il en va de même pour les rayons X et les rayons gamma car leur énergie est également délivrée sous forme d'électrons secondaires rapides. Les particules chargées lourdes ont un transfert d'énergie linéaire beaucoup plus élevé et l'équivalent de dose est supérieur à la dose absorbée, (Q est d'environ 20 pour les particules alpha d'énergies typiques). Étant donné que les neutrons fournissent la majeure partie de leur énergie sous la forme de particules chargées lourdes, leur facteur de qualité effectif est également considérablement supérieur à l'unité et varie considérablement avec l'énergie des neutrons.

Les unités utilisées pour l'équivalent de dose H dépendent des unités correspondantes de dose absorbée D . Si D est exprimé en radian, H est défini comme étant en unités de *rem*. Historiquement, le *rem* (ou *millirem*) était universellement utilisé pour quantifier l'équivalent de dose. En convention SI, D est plutôt exprimé en grays, et une unité correspondante d'équivalent de dose appelée *sievert* (Sv) est maintenant utilisée. Par exemple, une dose absorbée de 2 Gy délivrée par rayonnement avec $Q = 5$ se traduira par un équivalent de dose de 10 Sv. Comme 1 Gy = 100 rad, les unités sont liées par :

$$\boxed{1 \text{ Sv} = 100 \text{ rem}} \quad (26)$$

6.4 Conversion fluence-dose

Pour des rayonnements relativement pénétrants tels que les neutrons et les rayons gamma, il est parfois commode de pouvoir estimer l'équivalent de dose dont une personne est exposée qui résulte d'une fluence donnée du rayonnement. La fluence est définie comme $\Phi = dN/da$ où dN représente le nombre différentiel des photons du rayonnement gamma ou de neutrons qui sont incidents à une sphère avec une section transversale différentielle da . En termes de flux neutronique, la fluence est l'intégrale temporelle du flux sur la durée d'exposition. Pour un faisceau monodirectionnel, la fluence est simplement le nombre de photons ou de neutrons par unité de surface, où la surface est perpendiculaire à la direction du faisceau. Le nombre de comptages enregistrés à partir des détecteurs de rayonnement est plus facilement interprété en termes de nombre d'interactions induites par le rayonnement dans le détecteur, et est donc plus étroitement liée à la fluence qu'à la dose. Ainsi, une conversion entre la fluence et la dose est généralement effectuée lors de l'interprétation des mesures effectuées avec des détecteurs en mode impulsionnel pour déterminer la dose. Étant donné que l'atténuation des rayons gamma ou des neutrons rapides dans l'air est négligeable sur de petites distances, la fluence directe (ou non diffusée) peut souvent être estimée à partir de sources ponctuelles en supposant que $\Phi = N/4\pi d^2$, où N est le nombre de photons ou neutrons émis par la source et d la distance de la source.

La conversion de la fluence en dose doit tenir compte des probabilités dépendantes de l'énergie de produire des particules secondaires ionisantes à partir d'interactions de rayons gamma et/ou de neutrons, de l'énergie cinétique déposée par ces particules et du facteur de qualité Q qui devrait être appliqué pour convertir l'énergie déposée par unité de masse en équivalent de dose. En supposant un modèle physique pour le corps humain, les codes de transport des rayons gamma et des neutrons peuvent être appliqués pour calculer le dépôt d'énergie et la dose pour les différents tissus et systèmes d'organes. Ces calculs seront sensibles à la direction d'incidence du rayonnement en raison des effets d'auto-blindage et d'atténuation dans le corps. En tenant compte de la radiosensibilité différente des différents organes et tissus à travers les facteurs de pondération tissulaire recommandés par la CIPR (voir la section suivante), les composants de dose individuels peuvent ensuite être combinés en une dose efficace E pour représenter une estimation de l'effet biologique global d'une exposition uniforme du corps entier à la fluence supposée. La dose efficace s'écrit alors :

$$H_E = h_E \Phi \quad (27)$$

Où h_E est le coefficient de conversion de la fluence en dose efficace. Ces coefficients sont multipliés par la fluence dépendante de l'énergie incidente sur le corps et intégrés sur l'énergie pour obtenir la dose efficace.

6.5 Unités de dose en International Commission on Radiation Protection (ICRP)

La quantité de dose équivalente dans un organe ou un tissu T due au rayonnement R qui est représentée par le symbole $H_{T,R}$ est obtenue à partir de la dose absorbée $D_{T,R}$ moyennée sur un tissu ou un organe et multipliée par un facteur de pondération du rayonnement w_R qui tient compte des différents effets biologiques de divers rayonnements :

$$H_{T,R} = w_R \cdot D_{T,R} \quad (28)$$

La dose équivalente, contrairement à l'équivalent de dose ($H = DQ$), n'est pas une quantité ponctuelle mais une moyenne sur un tissu ou un organe. Les valeurs de w_R sont déterminées par l'énergie du rayonnement incident. Les unités de dose équivalente sont les sieverts lorsque la dose absorbée est exprimée en grays. S'il y a un mélange de rayonnements, alors la dose équivalente totale H_T est donnée par une somme pour tous les types de rayonnement :

$$H_T = \sum_R H_{T,R} = \sum_R w_R \cdot D_{T,R} \quad (29)$$

La *dose effective* E est définie comme la somme sur tous les tissus ou organes avec les *facteurs de pondération tissulaire* w_T fournis par l'ICRP :

$$E = \sum_T w_T \cdot H_T \quad (30)$$

6. 6 Quantités de dose opérationnelle

Comme la dose efficace E n'est pas une grandeur mesurable, l'ICRU a défini des *grandeurs opérationnelles* à utiliser dans les mesures pratiques. Les grandeurs opérationnelles sont définies séparément pour les rayonnements fortement pénétrants et pour les rayonnements faiblement pénétrants. Les rayonnements faiblement pénétrant sont, par exemple, le rayonnement bêta avec une énergie inférieure à 2 MeV ou les rayonnements photoniques avec des énergies inférieures à 15 KeV. Les rayonnements fortement pénétrants sont des photons à énergies élevées et des neutrons à toutes énergies. Les grandeurs opérationnelles sont basées sur l'équivalent de dose en un point du corps. À la mesure qu'ils dépendent du type et de l'énergie du rayonnement en ce point, ils peuvent être calculés à l'aide de la fluence de particules dépendante de l'énergie en ce point.

Il convient de souligner que l'estimation de l'effet d'une exposition donnée aux rayonnements ionisants est par nature une science inexacte. Les effets biologiques ne sont pas des grandeurs physiques absolues mesurables avec une grande précision. Les concepts de dose efficace ou d'équivalent de dose efficace ne visent qu'à fournir des indications pour estimer les effets potentiels d'une exposition donnée aux rayonnements et ne doivent pas être considérées comme des quantités très précises ou exactement reproductibles.

Chapitre 03

Propriétés générales des détecteurs de rayonnements

1. Introduction

Avant d'entamer l'étude des différents types de détecteurs de rayonnement, nous décrivons d'abord certaines propriétés générales qui s'appliquent à tous les types. Certaines définitions de base des propriétés des détecteurs, telles que l'efficacité et la résolution énergétique, ainsi que certains modes généraux de fonctionnement et méthodes d'enregistrement des données qui seront utiles pour catégoriser les applications des détecteurs seront inclus.

2. Modèle simplifié d'un détecteur

Nous débuterons avec un détecteur hypothétique soumis à un certain type d'irradiation. L'attention se porte d'abord sur l'interaction d'une seule particule ou quantum de rayonnement dans le détecteur, qui pourrait, par exemple, être une seule particule alpha ou un photon gamma. Pour que le détecteur réponde, le rayonnement doit subir une interaction à travers l'un des mécanismes vus au chapitre (02). Le temps d'interaction ou d'arrêt est très faible (typiquement quelques nanosecondes dans les gaz ou quelques picosecondes dans les solides). Dans la plupart des situations pratiques, ces temps sont si courts que le dépôt de l'énergie de rayonnement peut être considéré comme instantané.

Le résultat de l'interaction de rayonnement dans une large catégorie de détecteurs est l'apparition d'une quantité donnée de charge électrique dans le volume actif du détecteur. Le modèle simplifié du détecteur suppose donc qu'une charge Q apparaît dans le détecteur à l'instant $t = 0$ résultant de l'interaction d'une seule particule ou quantum de rayonnement. Ensuite, cette charge doit être collectée pour former un signal électrique de base. Typiquement, la collecte de la charge est accomplie par l'imposition d'un champ électrique à l'intérieur du détecteur, ce qui fait que les charges positives et négatives créées par le rayonnement circulent dans des directions opposées. Le temps nécessaire pour collecter complètement la charge varie considérablement d'un détecteur à l'autre. Par exemple, dans les chambres à ions (détecteur à gaz), le temps de collecte peut atteindre quelques millisecondes, alors que dans les détecteurs à diodes semi-conductrices, le temps est de quelques nanosecondes. Ces temps reflètent à la fois la mobilité des porteurs de charge dans le volume actif du détecteur et la distance moyenne qui doit être parcourue avant l'arrivée aux électrodes de collecte. Nous commençons donc avec un modèle de détecteur prototype dont la réponse à une seule particule ou quantum de rayonnement sera un courant qui circule pendant un temps égal au temps de collecte de la charge. La figure (01) illustre un exemple de la dépendance temporelle que le courant du détecteur pourrait assumer, où t_c représente le temps de collecte de charge.

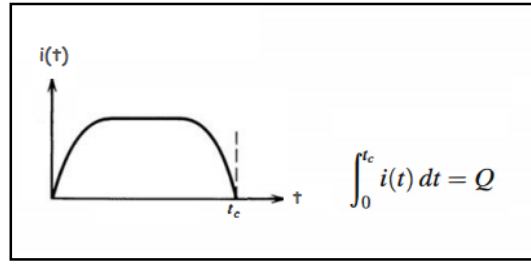


Fig. 01 Courant du détecteur circulant pendant un temps égal au temps de collecte de charge t_c

L'intégrale du courant durant la durée de collecte doit être simplement égale à Q , la quantité totale de charge générée dans cette interaction spécifique. Dans toute situation réelle, de nombreux quanta de rayonnement interagiront sur une période de temps. Si le taux d'irradiation est élevé, des situations peuvent survenir dans lesquelles un courant circule dans le détecteur à partir de plus d'une interaction à un instant donné. Finalement, nous supposons que le taux est suffisamment faible pour que chaque interaction individuelle donne lieu à un courant qui se distingue de tous les autres. L'amplitude et la durée de chaque impulsion de courant peuvent varier en fonction du type d'interaction, et un schéma du courant instantané circulant dans le détecteur peut alors apparaître comme indiqué dans la figure (02).

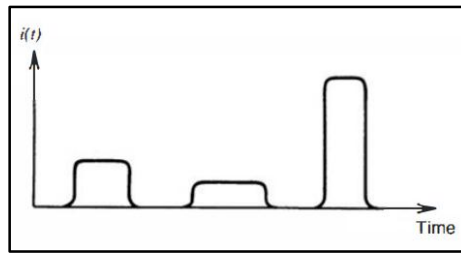


Fig. 02 Courant instantané circulant dans le détecteur (Faible taux d'irradiation du détecteur)

Il est important de rappeler que, l'arrivée des quanta de rayonnement étant un phénomène aléatoire régi par la statistique de Poisson, les intervalles de temps entre les impulsions de courant successives sont également distribués de manière aléatoire.

3. Modes de fonctionnement d'un détecteur

Nous pouvons maintenant introduire une distinction fondamentale entre trois modes généraux de fonctionnement des détecteurs de rayonnement. Les trois modes sont appelés mode impulsion, mode courant et mode moyenne carrée de voltage (mode MSV abrégé). Le mode pulsé est de loin le plus couramment appliqué, mais le mode courant trouve également de nombreuses applications. Le mode MSV est limité à certaines applications spécialisées qui utilisent ses caractéristiques uniques. Bien que les trois modes soient fonctionnellement distincts, ils sont interdépendants par leur

dépendance commune à la séquence d'impulsions de courant qui sont la sortie du modèle de détecteur simplifié.

En fonctionnement en mode impulsionnel, l'instrumentation de mesure est conçue pour enregistrer chaque quantum individuel de rayonnement qui interagit dans le détecteur. Dans la plupart des applications courantes, l'intégrale temporelle de chaque impulsion de courant, ou la charge totale Q , est enregistrée depuis l'énergie déposée dans le détecteur est directement liée à Q . Tous les détecteurs utilisés pour mesurer l'énergie des quanta de rayonnement individuels doivent fonctionner en mode impulsionnel. De telles applications sont classées comme spectroscopie de rayonnement.

Une approche additionnelle plus simple peut répondre aux besoins de la mesure, est que toutes les impulsions au-dessus d'un seuil bas sont enregistrées à partir du détecteur, quelle que soit la valeur de la charge Q . Cette approche est souvent appelée comptage d'impulsions. Il peut être utile dans de nombreuses applications dans lesquelles seule l'intensité du rayonnement est intéressante, plutôt que la distribution d'énergie incidente du rayonnement.

À des taux d'événements très élevés, le fonctionnement en mode pulsé devient peu pratique, voire impossible. Le temps entre des événements adjacents peut devenir trop court pour effectuer une analyse adéquate, où les impulsions de courant d'événements successifs peuvent se chevaucher dans le temps. Dans de tels cas, on peut revenir à des techniques de mesure alternatives qui répondent à la moyenne temporelle prise sur de nombreux événements individuels. Cette approche conduit aux deux modes de fonctionnement restants : le mode courant et le mode MSV.

3.1 Mode courant

Dans la figure (03), un appareil de mesure de courant (un ampèremètre ou, plus pratiquement, un pico-ampèremètre) est connecté aux bornes de sortie d'un détecteur de rayonnement.

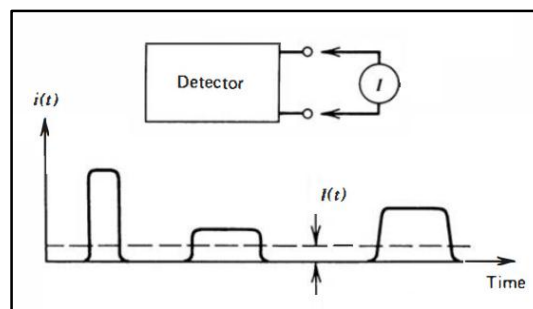


Fig. 03 Appareil de mesure de courant connecté aux bornes de sortie d'un détecteur de rayonnement.
Réponse d'un détecteur fonctionnant en mode courant

Si nous supposons que l'appareil de mesure a un temps de réponse fixe T , alors le signal enregistré à partir d'une séquence d'événements sera un courant dépendant du temps, donné par :

$$I(t) = \frac{1}{T} \int_{t-T}^t i(t') dt' \quad (01)$$

Étant donné que le temps de réponse T est généralement long par rapport au temps moyen entre les impulsions de courant individuelles du détecteur, l'effet est de moyenner de nombreuses fluctuations dans les intervalles entre les interactions de rayonnement individuelles et d'enregistrer un courant moyen qui dépend du produit du taux d'interaction et la charge moyenne par interaction. En mode courant, cette moyenne temporelle des courants individuels sert de signal de base qui est enregistré.

À tout instant, cependant, il existe une incertitude statistique dans ce signal en raison des fluctuations aléatoires du temps d'arrivée de l'événement. À bien des égards, le temps d'intégration T est analogue au temps de mesure. Ainsi, le choix d'un grand T minimisera les fluctuations statistiques du signal mais ralentira également la réponse aux changements rapides du taux ou nature des interactions de rayonnement. Le courant moyen est donné par le produit du taux d'événement et de la charge moyenne produite par événement.

$$I_0 = rQ = r \frac{E}{W} q \quad (02)$$

Avec :

- r : Taux d'événements par seconde ;
- Q : Charge produite par événement ;
- E : Energie moyenne déposée par événement ;
- e : Charge de l'électron = $1.6 \cdot 10^{-19}$ C;
- w : Energie moyenne requise pour produire une paire chargée.

Pour une irradiation en régime permanent du détecteur, ce courant moyen peut être également réécrit comme la somme d'un courant constant I_0 plus une composante fluctuante en fonction du temps $\sigma_i(t)$ comme l'illustre la figure (04) :

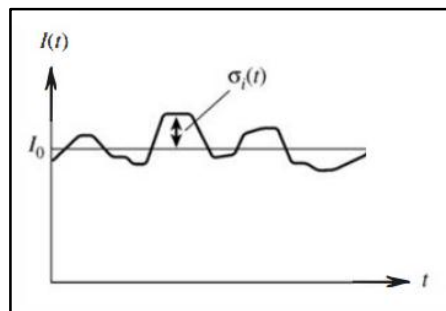


Fig. 04 Réponse à une irradiation en régime permanent du détecteur

Où, $\sigma_i(t)$ est une variable aléatoire dépendant du temps qui se produit en conséquence de la nature aléatoire des événements de rayonnement interagissant dans le détecteur.

Une mesure statistique de cette composante aléatoire est la variance ou la valeur quadratique moyenne, définie comme la moyenne temporelle du carré de la différence entre le courant fluctuant $I(t)$ et le courant moyen I_0 . Cette valeur quadratique moyenne est donnée par :

$$\overline{\sigma_I^2(t)} = \frac{1}{T} \int_{t-T}^t [I(t') - I_0]^2 dt' = \frac{1}{T} \int_{t-T}^t \sigma_i^2(t') dt' \quad (03)$$

Et l'écart type est :

$$\overline{\sigma_I(t)} = \sqrt{\overline{\sigma_I^2(t)}} \quad (04)$$

Rappelons que les statistiques de Poisson indiquent que l'écart type du nombre d'événements enregistrés n sur une période d'observation donnée devrait être :

$$\sigma_n = \sqrt{n} \quad (05)$$

Par conséquent, l'écart type du nombre d'événements se produisant à une vitesse r dans un temps de mesure effectif T est simplement :

$$\sigma_n = \sqrt{rT} \quad (06)$$

Si chaque impulsion contribue avec la même charge, l'écart type fractionnaire du signal mesuré dû aux fluctuations aléatoires du temps d'arrivée de l'impulsion est donné par :

$$\frac{\overline{\sigma_I(t)}}{I_0} = \frac{\sigma_n}{n} = \frac{1}{\sqrt{rT}} \quad (07)$$

Ici $\overline{\sigma_I(t)}$ est la moyenne temporelle de l'écart type du courant mesuré, T est le temps de réponse du pico mètre et I_0 est le courant moyen lu sur le compteur. Ce résultat est utile pour estimer l'incertitude associée à la mesure un mode courant donné.

Il convient de noter que, dans la dérivation de la dernière équation, la charge (Q) produite dans chaque événement est supposée constante. Par conséquent, le résultat ne tient compte que des fluctuations aléatoires du temps d'arrivée des impulsions, mais pas des fluctuations de l'amplitude des impulsions.

Dans certaines applications, cependant, cette deuxième source de variance du signal est faible par rapport à la première, et le caractère général des résultats donnés reste valable.

3.2 Mode de tension quadratique moyenne

L'étude des propriétés statistiques du signal en mode courant nous mène au mode de fonctionnement général suivant : le mode tension quadratique moyenne (MSV). Supposons que nous envoyons le signal de courant à travers un élément de circuit qui bloque le courant moyen I_0 et ne laisse passer que la composante fluctuante $\sigma_i(t)$. En fournissant des éléments de traitement de signal supplémentaires, nous calculons maintenant la moyenne temporelle de l'amplitude au carré de $\sigma_i(t)$. Les étapes de traitement sont illustrées dans la figure (05) :

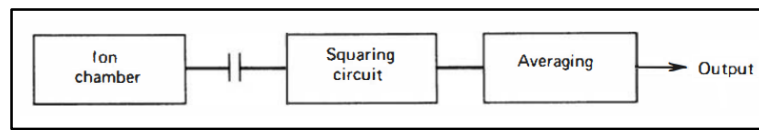


Fig. 05 Schéma bloc des étapes de traitement permettant le calcul de $\overline{\sigma_I^2(t)}$

Le résultat correspond à la quantité $\overline{\sigma_I^2(t)}$ définie précédemment. En combinant les équations (02) et (07), nous prévoyons que l'amplitude du signal dérivé de cette manière sera :

$$\boxed{\overline{\sigma_I^2(t)} = \frac{rQ^2}{T}} \quad (08)$$

On voit que le signal quadratique moyen est directement proportionnel au taux d'événement r et, plus significativement, proportionnel au carré de la charge Q produite à chaque événement.

Le mode de fonctionnement MSV est particulièrement utile pour effectuer des mesures dans des environnements à rayonnement mixte lorsque la charge produite par un type de rayonnement est très différente de celle du second type. Si le fonctionnement en mode courant simple est choisi, le courant mesuré reflétera linéairement les charges apportées par chaque type. En mode MSV, cependant, le signal dérivé est proportionnel au carré de la charge par événement. Ce mode de fonctionnement va donc pondérer davantage la réponse du détecteur en faveur du type de rayonnement donnant la charge moyenne la plus élevée par événement. Comme exemple d'application utile du mode MSV, il est utilisé avec des détecteurs de neutrons dans l'instrumentation du réacteur pour améliorer le signal neutronique par rapport à la réponse due aux événements des rayons gamma de plus faible amplitude.

3.3 Mode d'impulsion

En examinant diverses applications des détecteurs de rayonnement, nous constatons que le fonctionnement en mode courant est utilisé avec de nombreux détecteurs lorsque les taux d'événements sont très élevés. Les détecteurs appliqués à la dosimétrie des rayonnements fonctionnent également normalement en mode courant. Le mode MSV est utile pour améliorer la réponse relative aux événements de grande amplitude et trouve une large application dans l'instrumentation des réacteurs. La plupart des applications, cependant, sont mieux servies en préservant les informations sur l'amplitude et le moment des événements individuels que seul le mode pulsé peut fournir. Par conséquent, le reste de ce chapitre traitera divers aspects du fonctionnement en mode impulsionnel.

La nature de l'impulsion du signal produite à partir d'un événement unique dépend des caractéristiques d'entrée du circuit auquel le détecteur est connecté (généralement un préamplificateur). Le circuit équivalent peut être souvent représenté comme l'indique la figure (06).

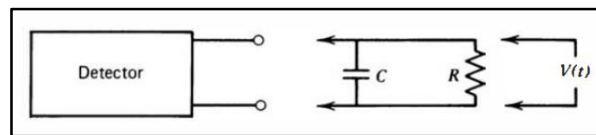


Fig. 06 Schéma bloc d'un circuit (détecteur connecté à un préamplificateur) fonctionnant en mode impulsionnel

R représente la résistance d'entrée du circuit et C représente la capacité équivalente du détecteur lui-même et du circuit de mesure. Si un préamplificateur est connecté au détecteur, alors (R) est sa résistance d'entrée et (C) est la capacité équivalente du détecteur, le câble utilisé pour connecter le détecteur au préamplificateur et la capacité d'entrée du préamplificateur lui-même. Dans la plupart des cas, la tension $V(t)$ aux bornes de la résistance de charge est la tension de signal fondamentale sur laquelle le fonctionnement en mode impulsionnel est basé. Deux états extrêmes distincts de fonctionnement peuvent être identifiés qui dépendent de la valeur relative de la constante de temps du circuit de mesure. À partir d'une simple analyse du circuit, cette constante de temps est donnée par le produit de R et C , ou $\tau = RC$.

3.3.1 Cas où RC est très faible ($\tau \ll t_c$)

Dans ce cas extrême, la constante de temps τ du circuit externe est faible par rapport au temps de collecte de charge t_c , de sorte que le courant traversant la résistance de charge R est essentiellement égal à la valeur instantanée du courant circulant dans le détecteur. La tension de signal $v(t)$ produite dans ces conditions a une forme presque identique à $i(t)$, courant produit dans le détecteur comme l'illustre la figure (07 b) suivante. Les détecteurs de rayonnement fonctionnent parfois dans ces

conditions lorsque les taux d'événements sont élevés ou les informations temporelles sont plus importantes que les informations énergétiques.

3.3.2 Cas où RC est très grande ($\tau \gg t_c$)

Il est généralement plus commode de faire fonctionner les détecteurs à l'état dans lequel la constante de temps du circuit externe est beaucoup plus grande que le temps de collecte de charge du détecteur. Dans ce cas, un très faible courant circulera dans la résistance de charge pendant le temps de collecte de charge et le courant du détecteur est momentanément intégré sur la capacité. Si nous supposons que le temps entre les impulsions est suffisamment grand, la capacité se déchargera alors à travers la résistance, ramenant la tension aux bornes de la résistance de charge à zéro. La tension de signal correspondante $v(t)$ est illustrée sur la figure (07 c).

Étant donné que ce dernier cas est de loin le moyen le plus courant de fonctionnement du type impulsionnel des détecteurs, il est important de tirer quelques conclusions générales. Tout d'abord, le temps nécessaire pour que l'impulsion du signal atteigne sa valeur maximale est déterminée par le temps de collecte de charges dans le détecteur lui-même. Aucune propriété du circuit externe ou de charge n'influence le temps de montée des impulsions. D'autre part, le temps de décroissance des impulsions, ou le temps nécessaire pour ramener la tension à zéro, n'est déterminé que par la constante de temps du circuit de charge. La conclusion selon laquelle le front montant dépend du détecteur et le front descendant dépend du circuit, est une généralité qui sera valable pour une grande variété de détecteurs de rayonnement fonctionnant dans les conditions dans lesquelles $RC \gg t_c$.

Deuxièmement, l'amplitude d'une impulsion du signal représentée par $V_{\max}$ sur la figure 06.c est déterminée tout simplement par le rapport de la charge totale Q créée dans le détecteur lors d'une interaction de rayonnement divisée par la capacité C du circuit de charge équivalent. Étant donné que cette capacité est normalement fixe, l'amplitude de l'impulsion du signal est directement proportionnelle à la charge correspondante générée dans le détecteur et est donnée par la simple expression :

$$V_{\max} = Q/C \quad (09)$$

Ainsi, la sortie d'un détecteur fonctionnant en mode impulsionnel consiste normalement en une séquence d'impulsions de signal individuelles, chacune représentant les résultats de l'interaction d'un seul quantum de rayonnement à l'intérieur du détecteur. Une mesure du taux auquel de telles impulsions se produisent donnera le taux correspondant d'interactions de rayonnement dans le détecteur. De plus, l'amplitude de chaque impulsion individuelle reflète la quantité de charge générée en raison de chaque interaction individuelle. Nous verrons qu'une méthode d'analyse très courante consiste à enregistrer la distribution de ces amplitudes à partir de laquelle certaines informations

peuvent souvent être déduites sur le rayonnement incident. Un exemple est que cet ensemble de conditions dans lesquelles la charge Q est directement proportionnelle à l'énergie du quantum incident de rayonnement. Ensuite, une distribution enregistrée des amplitudes d'impulsion reflétera la distribution correspondante en énergie du rayonnement incident.

D'après l'équation précédente, la proportionnalité entre $V_{\max}$ et Q n'est valable que si la capacité C reste constante. Dans la plupart des détecteurs, la capacité inhérente est définie par sa taille et sa forme, et l'hypothèse de constance est pleinement justifiée. Dans d'autres types (notamment le détecteur à diode semi-conductrice), la capacité peut changer avec les variations des paramètres de fonctionnement normaux. Dans de tels cas, des impulsions de tension d'amplitude différente peuvent résulter d'événements avec le même Q . Afin de préserver les informations de base portées par l'amplitude de Q , un type de circuit de préamplificateur connu sous le nom de configuration sensible à la charge est devenu largement utilisé. Ce type de circuit utilise la rétroaction pour éliminer en grande partie la dépendance de l'amplitude de sortie de la valeur de C et restaure la proportionnalité à la charge Q même dans les cas où C peut changer.

Pour cette configuration de préamplificateur, la simple représentation RC représentée dans la figure 06 n'est plus exacte, mais le principe de collecte de l'impulsion de courant à travers une capacité qui se décharge à travers une résistance reste valable.

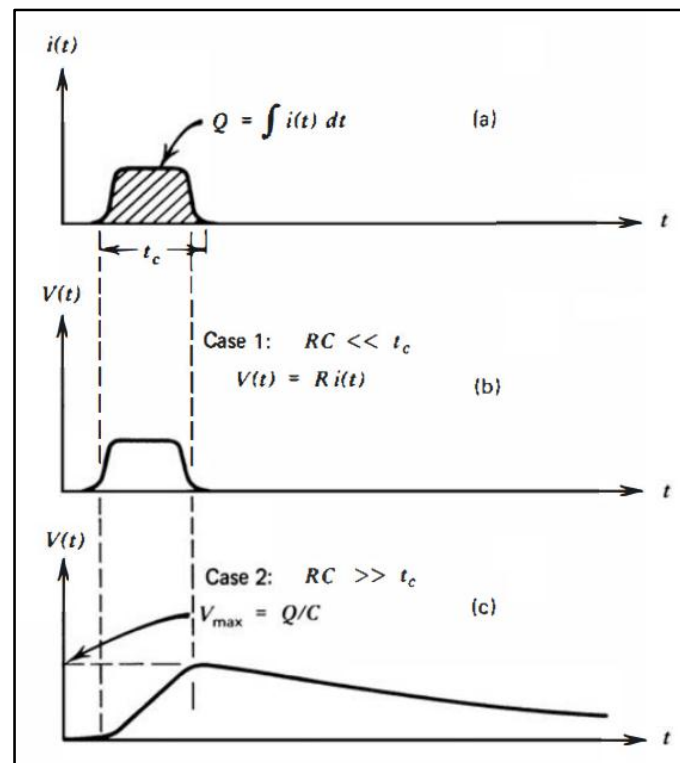


Fig. 07 (a) Courant de sortie d'un détecteur ;
 (b) Tension $v(t)$, cas d'un circuit de charge à constante de temps très petite ;
 (c) Tension $v(t)$, cas d'un circuit de charge à constante de temps très grande.

Le fonctionnement en mode impulsionnel est le choix le plus courant pour la plupart des applications de détection de rayonnement en raison de plusieurs avantages inhérents par rapport au mode courant. Tout d'abord, la sensibilité qui peut être obtenue est souvent bien supérieure à celle obtenue en mode courant ou MSV, car chaque quantum individuel de rayonnement peut être détecté comme une impulsion distincte. Les limites inférieures de détectabilité sont alors normalement fixées par les niveaux de rayonnement. En mode courant, le courant minimum détectable peut représenter un taux d'interaction moyen dans le détecteur plusieurs fois supérieur. Le deuxième avantage, le plus important, est que chaque amplitude d'impulsion transporte des informations qui sont souvent une partie utile ou même nécessaire d'une application particulière.

Dans le fonctionnement en modes courant et MSV, les informations sur les amplitudes d'impulsions individuelles sont perdues et toutes les interactions, quelle que soit l'amplitude, contribuent au courant mesuré moyen. En raison de ces avantages inhérents du mode impulsionnel, l'accent est mis dans l'instrumentation nucléaire en grande partie sur les circuits impulsionnels et les techniques de traitement des impulsions.

4. Spectres des hauteurs des impulsions

Lors du fonctionnement d'un détecteur de rayonnement en mode impulsionnel, chaque amplitude d'impulsion individuelle contient des informations importantes concernant la charge générée par cette interaction de rayonnement dans le détecteur. Si nous examinons un grand nombre de telles impulsions, leurs amplitudes ne seront pas toutes les mêmes. Les variations peuvent être dues soit à la différence entre les énergies des rayonnements, soit à des fluctuations de la réponse inhérente du détecteur aux rayonnements mono-énergétiques. La distribution de l'amplitude des impulsions est une propriété fondamentale de la sortie du détecteur qui est couramment utilisée pour déduire des informations sur le rayonnement incident ou le fonctionnement du détecteur lui-même.

La manière la plus courante d'afficher les informations sur l'amplitude des impulsions consiste à utiliser la distribution différentielle de la hauteur des impulsions. La figure (08 a) donne une distribution quelconque, l'axe des abscisses est une échelle d'amplitude d'impulsion linéaire allant de zéro à une valeur supérieure à toute amplitude d'impulsion observée à partir de la source. L'axe des ordonnées est le nombre différentiel dN d'impulsions observées avec une amplitude appartenant à l'incrément d'amplitude différentielle dH , divisé par cet incrément, soit dN / dH . L'axe horizontal a alors les unités d'amplitude d'impulsion (volts), tandis que l'axe vertical a les unités de l'inverse d'amplitude (volts^{-1}). Le nombre d'impulsions dont l'amplitude se situe entre deux valeurs spécifiques, H_1 et H_2 , peut être obtenu en intégrant la zone sous la distribution entre ces deux limites, comme indiqué par la zone hachurée sur la figure (08 a) :

$$\text{number of pulses with amplitude between } H_1 \text{ and } H_2 = \int_{H_1}^{H_2} \frac{dN}{dH} dH \quad (10)$$

Le nombre total d'impulsions N_0 représenté par la distribution peut être obtenu en intégrant l'aire sous tout le spectre :

$$N_0 = \int_0^{\infty} \frac{dN}{dH} dH \quad (11)$$

Des caractéristiques significatives sur la source des impulsions peuvent être tirées à partir de la forme de la distribution de hauteur d'impulsion différentielle. La hauteur d'impulsion maximale observée (H_5) est simplement le point le long de l'abscisse auquel la distribution tend vers zéro. Des pics dans la distribution, comme en H_4 , indiquent des amplitudes d'impulsions autour desquelles un grand nombre d'impulsions peuvent être trouvées. D'autre part, les vallées ou les points bas dans le spectre, comme à la hauteur d'impulsion H_3 , indique des valeurs de l'amplitude d'impulsion autour desquelles relativement peu d'impulsions se produisent. L'interprétation physique des spectres de hauteur d'impulsion différentiels implique toujours des zones sous le spectre entre deux limites données de hauteur d'impulsion. La valeur de l'ordonnée elle-même (dN / dH) n'a aucune signification physique tant qu'elle n'est pas multipliée par un incrément de l'abscisse H .

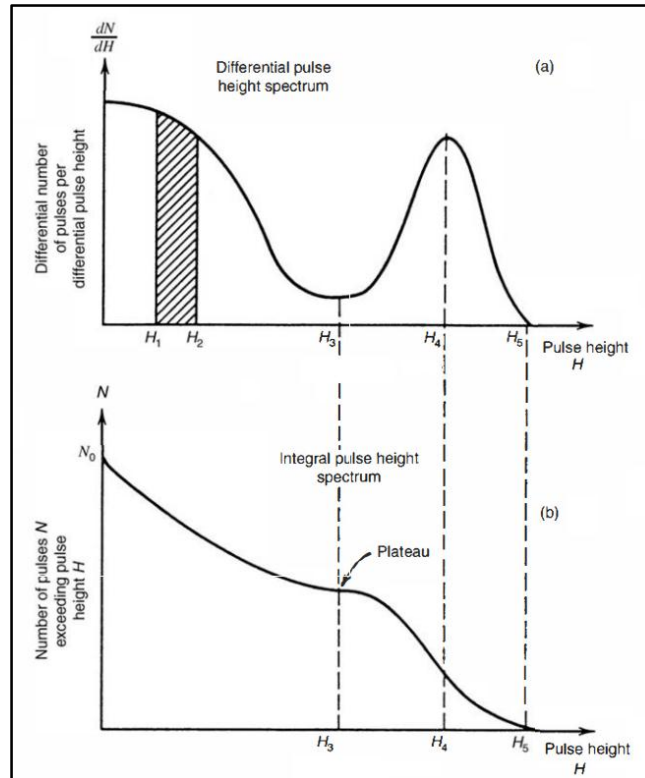


Fig. 08 Exemples de spectres de hauteur d'impulsion différentiel (a) et intégral (b) pour une source d'impulsions donnée

Une manière moins courante d'afficher les mêmes informations sur la distribution des amplitudes d'impulsion consiste à utiliser la distribution intégrale de la hauteur d'impulsion. La figure 08 (b) montre la distribution intégrale pour la même source d'impulsions affichée sous forme de spectre différentiel sur la figure 08 (a). L'axe des abscisses dans le cas de la distribution intégral est la même que pour la distribution différentielle.

L'axe des ordonnées représente dans ce cas le nombre d'impulsions dont l'amplitude dépasse celle d'une valeur donnée de l'abscisse H . L'ordonnée N doit être toujours une fonction monotone décroissante de H car de moins en moins d'impulsions se situent au-dessus d'une amplitude H qui peut croître à partir de zéro. Étant donné que toutes les impulsions ont une amplitude finie, la valeur du spectre intégral à $H = 0$ doit être le nombre total d'impulsions observées (N_0). La valeur de la distribution intégrale doit décroître à zéro à la hauteur d'impulsion maximale observée (H_5).

Les distributions différentielle et intégrale transmettent exactement la même information et l'une peut être dérivée de l'autre. L'amplitude de la distribution différentielle à toute hauteur d'impulsion H est donnée par la valeur absolue de la pente de la distribution intégrale à la même valeur. Lorsque des pics apparaissent dans la distribution différentielle, comme (H_4), des maxima locaux se produiront dans l'amplitude de la pente de la distribution intégrale. D'autre part, là où des minima apparaissent dans le spectre différentiel, comme H_3 , des régions d'amplitude minimale de la pente sont observées dans la distribution intégrale. Puisqu'il est plus facile d'afficher des différences subtiles en utilisant la distribution différentielle, elle est devenue le moyen prédominant d'afficher les informations de la distribution des hauteurs d'impulsion.

5. Courbes de comptage et plateaux

Lorsque des détecteurs de rayonnement fonctionnent en mode de comptage d'impulsions, une situation courante se produit souvent dans laquelle les impulsions du détecteur sont transmises à un dispositif de comptage avec un niveau de discrimination fixe. Les impulsions de signal doivent dépasser un niveau donné H_d pour être enregistrées par le circuit de comptage. Il est parfois possible de faire varier le niveau H_d au cours de la mesure pour renseigner sur la répartition en amplitude des impulsions. En supposant que H_d puisse varier de 0 à H_5 sur la figure 08, une série de mesures peut être effectuée dans laquelle le nombre d'impulsions N par unité de temps est mesuré lorsque H_d prend des valeurs entre 0 et H_5 . Cette série de mesures n'est qu'une détermination expérimentale de la distribution intégrale de la hauteur des impulsions, et les comptages mesurés doivent se situer directement sur la courbe illustrée par la figure (08 b).

Lors de la configuration d'une mesure de comptage d'impulsions, il est souvent souhaitable d'établir un point de fonctionnement qui fournira une stabilité maximale sur de grands intervalles de temps. Par exemple, de petites dérives de la valeur de H_d pourraient être attendues dans n'importe quelle

application réelle, et on aimerait établir des conditions dans lesquelles ces dérives auraient une influence minimale sur les comptages mesurés. Un tel point de fonctionnement stable peut être atteint à un point de discrimination fixé au niveau H_3 sur la figure 08. Étant donné que la pente de la distribution intégrale est minimale à ce point, de petites modifications du niveau de discrimination auront un impact minimal sur le nombre total d'impulsions enregistrées. En général, les régions de pente minimale sur la distribution intégrale sont appelées plateaux de comptage et représentent des zones de fonctionnement dans lesquelles une sensibilité minimale aux dérives du niveau de discrimination est atteinte. Il convient de noter que les plateaux dans le spectre intégral correspondent à des vallées dans la distribution différentielle.

Des plateaux dans les données de comptage peuvent également être observés avec une procédure différente. Pour un détecteur de rayonnement particulier, il est souvent possible de faire varier le gain ou l'amplification apportée à la charge produite dans les interactions de rayonnement. Cette variation pourrait être accomplie en faisant varier le facteur d'amplification d'un amplificateur linéaire entre le détecteur et le circuit de comptage, ou dans de nombreux cas plus directement en modifiant la tension appliquée au détecteur lui-même. La figure (09) montre la distribution différentielle de la hauteur des impulsions correspondant à trois valeurs différentes de la tension appliquée à la même source d'impulsions ce qui conduit à trois valeurs différentes du gain d'amplification. La valeur du gain peut être définie comme le rapport de l'amplitude de la tension pour un événement donné dans le détecteur à la même amplitude avant qu'un paramètre (tel que l'amplification ou la tension du détecteur) ne soit modifié. Le gain de la tension le plus élevé se traduira par la plus grande hauteur d'impulsion maximale, mais dans tous les cas, la zone sous la distribution différentielle sera une constante.

Dans l'exemple de la figure (09), aucun comptage ne sera enregistré pour un gain $G = 1$ car dans ces conditions toutes les impulsions seront inférieures à H_d . Des impulsions commenceront à être enregistrées quelque part entre un gain $G = 1$ et $G = 2$. Une expérience peut être réalisée dans laquelle le nombre d'impulsions enregistrées est mesuré en fonction du gain appliqué, parfois appelée courbe de comptage. Un tel tracé est également donné dans la figure 09 et ressemble à bien des égards à une distribution intégrale de hauteur d'impulsion. Nous avons maintenant une image miroir de distribution intégrale, car de petites valeurs du gain n'enregistreront aucune impulsion, tandis que de grandes valeurs entraîneront le comptage de presque toutes les impulsions. De nouveau, des plateaux peuvent être anticipés dans cette courbe de comptage pour des valeurs du gain dans lesquelles la hauteur d'impulsion de discrimination effective H_d passe par des minima dans la distribution de hauteur d'impulsion différentielle. Dans l'exemple illustré à la figure (09), la pente minimale de la

courbe de comptage doit correspondre à un gain d'environ 3, dans un tel cas le point de discrimination est proche du minimum de la vallée dans la distribution différentielle de hauteur d'impulsion.

Dans certains types de détecteurs de rayonnement, tels que les tubes Geiger-Mueller ou les compteurs à scintillation, le gain peut être commodément modifié en modifiant la tension appliquée au détecteur. Bien que le gain puisse ne pas changer linéairement avec la tension, les caractéristiques qualitatives de la courbe de comptage peuvent être tracées par une simple mesure du taux de comptage du détecteur en fonction de la tension. Afin de sélectionner un point de fonctionnement de stabilité maximale, des plateaux sont à nouveau recherchés dans la courbe de comptage qui en résulte, et la tension est souvent choisie pour se situer à un point de pente minimale sur cette courbe de comptage.

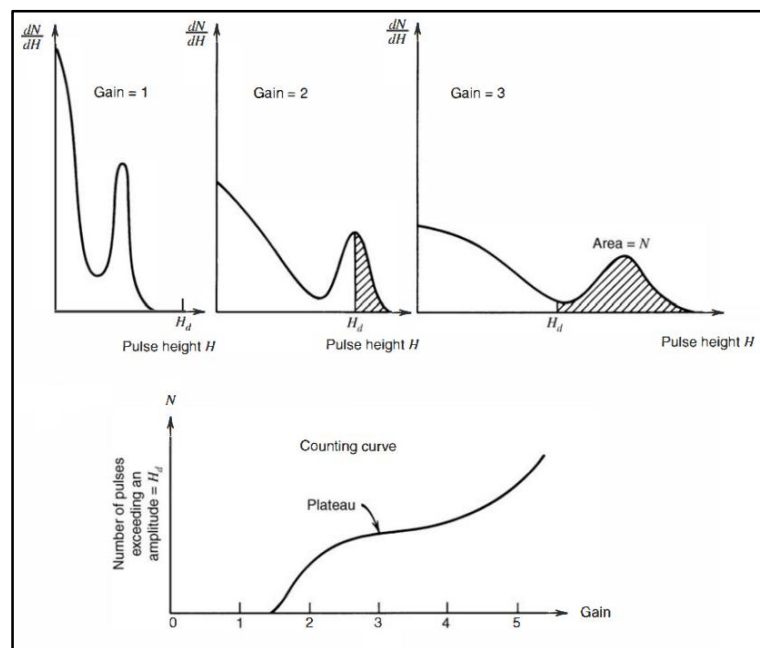


Fig. 09 Exemple de courbe de comptage générée en faisant varier le gain d'une source dans des conditions constantes. Les trois courbes en haut donnent le spectre différentiel de hauteur d'impulsion (courbe en bas)

6. Résolution en énergie

Dans de nombreuses applications des détecteurs de rayonnement, l'objectif est de mesurer la distribution d'énergie du rayonnement incident. Ces efforts sont classés sous le terme général de spectroscopie de rayonnement, par la suite nous donneront des exemples d'utilisation de détecteurs spécifiques pour la spectroscopie impliquant des particules alpha, des rayons gamma et d'autres types de rayonnement nucléaire. À ce stade, nous discutons de certaines propriétés générales des détecteurs lorsqu'ils sont appliqués à la spectroscopie de rayonnement et introduisons certaines définitions qui seront utiles dans ces discussions.

Une propriété importante d'un détecteur en spectroscopie de rayonnement peut être examinée en notant sa réponse à une source mono-énergétique de ce rayonnement. La figure (10) montre la

distribution différentielle de hauteur d'impulsion qui pourrait être produite par un détecteur dans ces conditions. Cette distribution est appelée fonction de réponse du détecteur pour l'énergie utilisée dans la détermination. La courbe avec "Bonne résolution" illustre une distribution possible autour d'une hauteur d'impulsion moyenne H_0 . La deuxième courbe, "Mauvaise résolution", illustre la réponse d'un détecteur aux performances inférieures. Si le même nombre d'impulsions est enregistré dans les deux cas, les aires sous chaque pic sont égales. Bien que les deux distributions soient centrées sur la même valeur moyenne H_0 , la largeur de la distribution dans le cas de mauvaise résolution est beaucoup plus grande. Cette largeur reflète le fait qu'une grande quantité de fluctuation a été enregistrée d'une impulsion à l'autre même si la même énergie a été déposée dans le détecteur pour chaque événement. Si la quantité de ces fluctuations est réduite, la largeur de la distribution correspondante deviendra également plus petite et le pic se rapprochera d'un pic aigu ou d'une fonction mathématique delta.

La capacité d'une mesure donnée à résoudre les détails précis de l'énergie incidente du rayonnement est évidemment améliorée à mesure que la largeur de la fonction de réponse, donnée à la figure (10) devient de plus en plus petite.

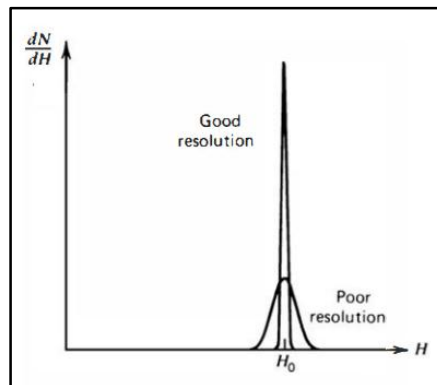


Fig. 10 Exemples de fonctions de réponse pour des détecteurs avec une bonne résolution relativement et une faible résolution relativement

Une définition formelle de la résolution en énergie d'un détecteur est présentée par la figure (11). La distribution de hauteur d'impulsion différentielle pour un détecteur hypothétique est représentée sous la même hypothèse que seul le rayonnement pour une seule énergie est enregistré. La pleine largeur à mi-hauteur (FWHM) est illustrée sur la figure (11) et est défini comme la largeur de la distribution à un niveau qui est juste la moitié de l'ordonnée maximale du pic. Cette définition suppose que tout fond ou continuum sur lequel le pic peut être superposé est négligeable ou a été soustrait. La résolution en énergie du détecteur est classiquement définie comme la FWHM divisée par l'emplacement du pic centré en H_0 . La résolution en énergie R est donc une fraction sans dimension, exprimée en pourcentage. Les détecteurs à diodes semi-conductrices utilisés en spectroscopie alpha peuvent avoir une résolution en énergie inférieure à 1 %, tandis que les détecteurs

à scintillation utilisés en spectroscopie gamma affichent normalement une résolution en énergie de l'ordre de 3 à 10 %. Il devrait être clair que plus la valeur de la résolution en énergie est petite, mieux le détecteur sera capable de distinguer deux rayonnements dont les énergies sont proches l'une de l'autre. Une règle empirique approximative est que l'on devrait être capable de résoudre deux énergies qui sont séparées par plus d'une valeur du détecteur FWHM.

Il existe un certain nombre de sources potentielles de fluctuation dans la réponse d'un détecteur donné qui entraînent une résolution en énergie imparfaite. Ceux-ci incluent toute dérive des caractéristiques de fonctionnement du détecteur au cours des mesures, les sources de bruit aléatoire dans le détecteur et le système d'instrumentation, et le bruit statistique résultant de la nature discrète du signal mesuré lui-même. La troisième source est en quelque sorte la plus importante car elle représente une quantité minimale irréductible de fluctuation qui sera toujours présente dans le signal du détecteur, quelle que soit la perfection du reste du système. Dans une large catégorie d'applications de détection, le bruit statistique représente la principale source de fluctuation du signal et fixe ainsi une limite importante sur les performances du détecteur.

Le bruit statistique provient du fait que la charge Q générée à l'intérieur du détecteur par un quantum de rayonnement n'est pas une variable continue mais représente plutôt un nombre discret de porteurs de charge. Par exemple, dans une chambre d'ionisation, les porteurs de charge sont les paires d'ions produites par le passage de la particule chargée à travers la chambre, alors que dans un compteur à scintillation, ce sont le nombre d'électrons collectés à partir de la photocathode du tube photomultiplicateur.

Dans tous les cas, le nombre de porteurs est discret et sujet à des fluctuations aléatoires d'un événement à l'autre même si exactement la même quantité d'énergie est déposée dans le détecteur.

Une estimation peut être faite de la quantité de fluctuation inhérente en supposant que la formation de chaque porteur de charge est un processus de Poisson. Sous cette hypothèse, si un nombre total N de porteurs de charge est généré en moyenne, on s'attendrait à ce qu'un écart type de $\sqrt{N}$ caractérise les fluctuations statistiques inhérentes à ce nombre. S'il s'agissait de la seule source de fluctuation du signal, la fonction de réponse, comme le montre la figure (11), devrait avoir une forme gaussienne, car N est généralement un grand nombre. Dans ce cas, la fonction gaussienne s'écrit :

$$G(H) = \frac{A}{\sigma\sqrt{2\pi}} \exp\left(-\frac{(H - H_0)^2}{2\sigma^2}\right) \quad (12)$$

Le paramètre de largeur σ détermine la FWHM de toute gaussienne par la relation $\text{FWHM} = 2.35\sigma$ (Les deux paramètres restants, H_0 et A , représentent respectivement le centre et l'aire).

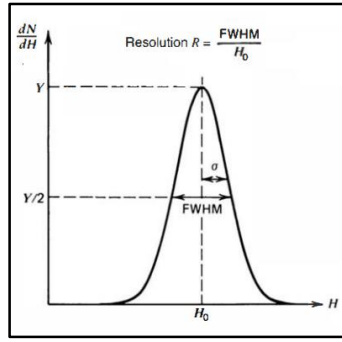


Fig. 11 Résolution du détecteur pour les pics dont la forme est gaussienne avec un écart type σ , la FWHM est donnée par $2,35 \sigma$

La réponse de nombreux détecteurs est approximativement linéaire, de sorte que l'amplitude moyenne des impulsions $H_0 = KN$, où K est une constante de proportionnalité. L'écart type σ du pic dans le spectre de hauteur d'impulsion est alors $\sigma = K\sqrt{N}$ et son FWHM est de $2.35K\sqrt{N}$. Nous calculerions alors une résolution limite R due uniquement aux fluctuations statistiques du nombre de porteurs de charge comme :

$$R|_{\text{Poisson limit}} \equiv \frac{\text{FWHM}}{H_0} = \frac{2.35K\sqrt{N}}{KN} = \frac{2.35}{\sqrt{N}} \quad (13)$$

Notons que cette valeur limite de R ne dépend que du nombre de porteurs de charge N , et la résolution s'améliore (R diminuera) à mesure que N augmente. De l'équation précédente on voit que pour atteindre une résolution en énergie meilleure que 1 %, il faut avoir N supérieur à 55 000. Un détecteur idéal aurait autant de porteurs de charge générés par événement que possible, de sorte que cette résolution limite soit à un pourcentage aussi faible que possible. Les détecteurs à semi-conducteurs sont caractérisés par le fait qu'un très grand nombre de porteurs de charge sont générés dans ces dispositifs par unité d'énergie perdue par le rayonnement incident.

Des mesures minutieuses de la résolution en énergie de certains types de détecteurs de rayonnement ont montré que les valeurs réalisables pour R peuvent être inférieures d'un facteur aussi grand que 3 ou 4 au minimum prédit par les arguments statistiques donnés ci-dessus. Ces résultats indiqueraient que les processus qui donnent lieu à la formation de chaque porteur de charge individuel ne sont pas indépendants, et donc le nombre total de porteurs de charge ne peut pas être décrit par de simples statistiques de Poisson. Le facteur de Fano a été introduit pour tenter de quantifier l'écart des fluctuations statistiques observées dans le nombre de porteurs de charge par rapport aux statistiques de Poisson pures et est défini comme :

$$F \equiv \frac{\text{observed variance in } N}{\text{Poisson predicted variance}(= N)} \quad (14)$$

Puisque la variance est donnée par σ^2 l'expression équivalente $R|_{\text{Statistical limit}}$ devient :

$$R|_{\text{Statistical limit}} = \frac{2.35 K \sqrt{N} \sqrt{F}}{KN} = 2.35 \sqrt{\frac{F}{N}} \quad (15)$$

Bien que le facteur de Fano soit sensiblement inférieur à l'unité pour les détecteurs à diode semi-conductrice et les compteurs proportionnels, d'autres types tels que de nombreux détecteurs à scintillation semblent montrer une résolution limite compatible avec les statistiques de Poisson et le facteur de Fano serait, dans ces cas, égal à l'unité.

Le fait que des facteurs de Fano inférieurs à un soient observés pour certains détecteurs peut être considéré comme une conséquence du processus de partition de l'énergie d'origine transportée par la particule incidente. Un simple argument suffit pour montrer que l'on s'attendrait à observer un facteur de Fano inférieur à l'unité dans le résultat d'expériences d'échantillonnage de nature générale dans lesquelles il existe une contrainte sur la quantité totale soumise à l'échantillonnage. Dans le cas d'une particule ionisante perdant son énergie dans le matériau du détecteur, un "échantillon" représente la quantité d'énergie qui entre dans la création d'un porteur de charge, et l'hypothèse est faite que cette énergie est sujette à des variations d'un échantillon à un autre. Cette variation reflète différentes quantités d'énergie qui sont perdues dans les processus thermiques lorsque des porteurs de charge individuels se forment le long de la trajectoire des particules.

En ajoutant la contrainte que la somme de toutes les pertes d'énergie doit être égale à l'énergie initiale des particules, il est alors prédit que le nombre de porteurs de charge produits fluctuera avec une variance inférieure à celle d'une distribution de Poisson.

Toute autre source de fluctuations dans la chaîne du signal se combinera avec les fluctuations statistiques inhérentes au détecteur pour donner la résolution énergétique globale du système de mesure. Il est parfois possible de mesurer la contribution à la FWHM globale due à un seul composant. Par exemple, si le détecteur est remplacé par un générateur d'impulsions stable, la réponse mesurée du reste du système montrera une fluctuation due principalement au bruit électronique. S'il existe plusieurs sources de fluctuation présentes et que chacune est symétrique et indépendante, la théorie statistique prédit que la fonction de réponse globale tendra toujours vers une forme gaussienne, même si les sources individuelles sont caractérisées par des distributions de forme différente. En conséquence, la fonction gaussienne donnée dans l'équation de $G(H)$ est largement utilisée pour représenter la fonction de réponse des systèmes de détection dans lesquels de nombreux facteurs différents peuvent contribuer à la résolution globale en énergie. Ensuite, le FWHM total sera la somme quadrature des valeurs FWHM pour chaque source individuelle de fluctuation :

$$\boxed{(\text{FWHM})_{\text{overall}}^2 = (\text{FWHM})_{\text{statistical}}^2 + (\text{FWHM})_{\text{noise}}^2 + (\text{FWHM})_{\text{drift}}^2 + \dots} \quad (16)$$

Chaque terme dans le deuxième membre de l'équation est le carré de la FWHM qui serait observée si toutes les autres sources de fluctuation étaient nulles.

7. Efficacité de détection

Tous les détecteurs de rayonnement donneront, en principe, lieu à une impulsion de sortie pour chaque quantum de rayonnement qui interagit dans son volume actif. Pour les rayonnements chargés primaires tels que les particules alpha ou bêta, l'interaction sous forme d'ionisation ou d'excitation aura lieu immédiatement dès l'entrée de la particule dans le volume actif. Après avoir parcouru une petite fraction de sa trajectoire, une particule typique formera suffisamment de paires d'ions le long de son trajet pour s'assurer que l'impulsion résultante est suffisamment grande pour être enregistrée. Ainsi, il est souvent facile d'organiser une situation dans laquelle un détecteur verra chaque particule alpha ou bêta qui entre dans son volume actif. Dans ces conditions, le détecteur est dit avoir une efficacité de comptage de 100%.

D'autre part, les rayonnements non chargés tels que les rayons gamma ou les neutrons doivent d'abord subir une interaction importante dans le détecteur avant que la détection ne soit possible. Étant donné que ces rayonnements peuvent parcourir de grandes distances entre les interactions, les détecteurs sont souvent moins efficaces à 100 %. Il devient alors nécessaire de disposer d'une valeur précise du rendement du détecteur afin de rapporter le nombre d'impulsions comptées au nombre de neutrons ou de photons incidents sur le détecteur.

Il convient de subdiviser les efficacités de comptage en deux classes : absolue et intrinsèque.

L'efficacité absolue est définie comme :

$$\boxed{\epsilon_{\text{abs}} = \frac{\text{number of pulses recorded}}{\text{number of radiation quanta emitted by source}}} \quad (17)$$

Elle dépend non seulement des propriétés du détecteur mais aussi des détails de la géométrie de comptage (principalement la distance entre la source et le détecteur). L'efficacité intrinsèque est définie comme :

$$\boxed{\epsilon_{\text{int}} = \frac{\text{number of pulses recorded}}{\text{number of radiation quanta incident on detector}}} \quad (18)$$

Elle n'inclut plus l'angle solide sous-tendu par le détecteur comme facteur implicite. Les deux efficacités sont simplement liées pour les sources isotropes par : $\epsilon_{\text{int}} = \epsilon_{\text{abs}} \cdot (4\pi/\Omega)$, où Ω est l'angle solide du détecteur vu depuis la position réelle de la source. Il est beaucoup plus pratique de tabuler des valeurs d'efficacités intrinsèques plutôt qu'absolues car la dépendance géométrique est beaucoup plus faible pour les premières. L'efficacité intrinsèque d'un détecteur dépend principalement du matériau du détecteur, de l'énergie de rayonnement et de l'épaisseur physique du détecteur dans la direction du rayonnement incident. Une légère dépendance à la distance entre la source et le détecteur est à considérer, parce que la longueur moyenne de la trajectoire du rayonnement à travers le détecteur changera quelque peu avec cette distance.

Les efficacités de comptage sont également catégorisées selon la nature de l'événement enregistré. Si nous considérons toutes les impulsions du détecteur, il convient alors d'utiliser les efficacités totales. Dans ce cas, toutes les interactions, quelle que soit leur faible énergie, sont supposés être comptés. En termes de distribution de hauteur d'impulsion différentielle hypothétique illustrée à la figure (11), la zone entière sous le spectre est une mesure du nombre de toutes les impulsions enregistrées, quelle que soit l'amplitude, et serait comptée dans la définition de l'efficacité totale. En pratique, tout système de mesure impose toujours une exigence selon laquelle les impulsions doivent être supérieures à un certain niveau de seuil fini, défini pour discriminer les très petites impulsions provenant de sources de bruit électronique. Ainsi, on ne peut s'approcher du rendement total théorique qu'en fixant ce niveau de seuil le plus bas possible. L'efficacité maximale, cependant, suppose que seuls ceux les interactions qui déposent toute l'énergie du rayonnement incident sont comptées. Dans une distribution de hauteur d'impulsion différentielle, ces événements à pleine énergie sont normalement mis en évidence par un pic qui apparaît à l'extrémité la plus élevée du spectre. Les événements qui ne déposent qu'une partie de l'énergie de rayonnement incident apparaîtront alors plus à gauche dans le spectre. Le nombre d'événements à pleine énergie peut être obtenu en intégrant simplement la surface totale sous le pic, qui est représentée par la zone hachurée sur la figure (12). Les efficacités totales et maximales (au pic) sont liées par le rapport r (*peak-to-total*) :

$$r = \frac{\epsilon_{\text{peak}}}{\epsilon_{\text{total}}} \quad (19)$$

Il est souvent préférable, d'un point de vue expérimental, de n'utiliser que les efficacités de pic, car le nombre d'événements à pleine énergie n'est pas sensible à certains effets perturbateurs tels que la diffusion par les objets environnants ou les bruits parasites.

Par conséquent, les valeurs de l'efficacité maximale peuvent être compilées et universellement appliquées à une grande variété de conditions de laboratoire, tandis que les valeurs d'efficacité totale peuvent être influencées par des conditions variables.

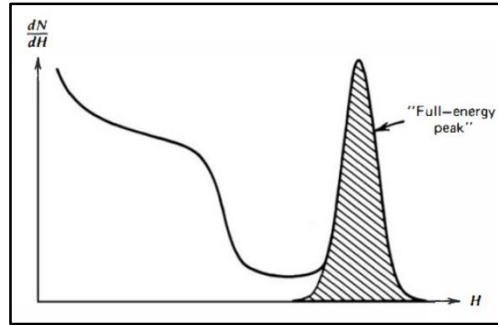


Fig. 12 Pic de pleine énergie dans un spectre différentiel de hauteur d'impulsion

Pour être complet, l'efficacité d'un détecteur doit être spécifiée selon les deux critères ci-dessus. Par exemple, le type d'efficacité le plus courant tabulé pour les détecteurs de rayons gamma est l'efficacité de crête (pic) intrinsèque.

Un détecteur d'efficacité connue peut être utilisé pour mesurer l'activité absolue d'une source radioactive. Par la suite nous supposons qu'un détecteur avec une efficacité de crête intrinsèque ϵ_{ip} a été utilisé pour enregistrer N événements sous le pic de pleine énergie dans le spectre du détecteur. Pour simplifier, nous supposons également que la source émet un rayonnement de manière isotrope et qu'aucune atténuation n'a lieu entre la source et le détecteur. A partir de la définition de l'efficacité de crête intrinsèque, le nombre de quanta de rayonnement S émis par la source pendant la période de mesure est alors donné par :

$$S = N \frac{4\pi}{\epsilon_{ip}\Omega} \quad (20)$$

Où Ω représente l'angle solide (en stéradians) sous-tendu par le détecteur à la position de la source. L'angle solide est défini par une intégrale sur la surface du détecteur qui fait face à la source, de la forme :

$$\Omega = \int_A \frac{\cos \alpha}{r^2} dA \quad (21)$$

Où r représente la distance entre la source et un élément de surface dA , et α est l'angle entre la normale à l'élément de surface et la direction de la source. Si le volume de la source est non

négligeable, alors une deuxième intégration doit être effectuée sur tous les éléments volumiques de la source. Pour le cas courant d'une source ponctuelle située le long de l'axe d'un détecteur cylindrique circulaire droit, Ω est donné par :

$$\Omega = 2\pi \left(1 - \frac{d}{\sqrt{d^2 + a^2}} \right) \quad (22)$$

Où la distance source-détecteur d et le rayon du détecteur a sont indiqués dans la figure (13) :

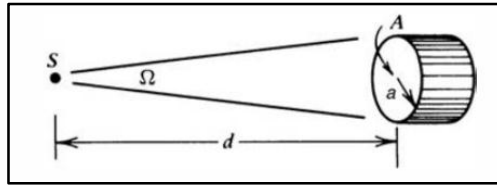


Fig. 13 Angle solide sous lequel est vu un détecteur depuis une source ponctuelle

Pour $d \gg a$, l'angle solide se réduit au rapport de la surface frontale du plan du détecteur A visible à la source au carré de la distance :

$$\Omega \cong \frac{A}{d^2} = \frac{\pi a^2}{d^2} \quad (23)$$

Un autre cas couramment rencontré, illustrée dans figure (14), implique une source de face circulaire uniforme émettant un rayonnement isotrope aligné avec un détecteur de face circulaire, tous deux positionnés perpendiculairement à un axe commun passant par leurs centres :

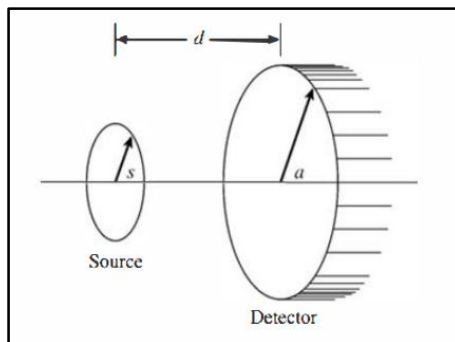


Fig. 14 Source uniforme (non ponctuelle) émettant un rayonnement isotrope vers un détecteur

En termes de dimensions indiquées sur la figure (14), on peut montrer que l'angle solide effectif moyenné sur la surface de la source est donné en résolvant l'intégrale suivante, cette dernière n'a pas de solution analytique, il ne peut être résolu qu'à l'aide de techniques numériques.

$$\Omega = \frac{4\pi a}{s} \int_0^\infty \frac{\exp(-dk) J_1(sk) J_1(ak)}{k} dk \quad (24)$$

Où $J_1(x)$ sont les fonctions de Bessel en x .

En termes de dimensions indiquées sur le croquis, on peut montrer que l'angle solide effectif moyenné sur la surface de la source est donné en résolvant une intégrale qui n'a pas de solution analytique, il ne peut donc être résolu qu'à l'aide de techniques numériques.

8. Temps mort

Dans presque tous les systèmes de détection, il y aura un minimum de temps qui doit séparer deux événements afin qu'ils soient enregistrés comme deux impulsions distinctes. Dans certains cas, le temps limite peut être défini par des processus dans le détecteur lui-même, et dans d'autres cas, la limite peut survenir dans l'électronique associée. Cette séparation temporelle minimale est généralement appelée temps mort du système de comptage. En raison de la nature aléatoire de la désintégration radioactive, il y a toujours une certaine probabilité qu'un véritable événement soit perdu parce qu'il se produit trop rapidement après un événement précédent. Ces "pertes de temps mort" peuvent devenir assez graves lorsque des taux de comptage élevés sont rencontrés, et toute mesure de comptage précise effectuée dans ces conditions doit inclure une certaine correction pour ces pertes. Dans cette section, nous aborderons quelques modèles de comportement de temps mort des systèmes de comptage, ainsi que deux méthodes expérimentales de détermination du temps mort du système.

8.1 Modèles de comportement du temps mort

Deux modèles de comportement du temps mort des systèmes de comptage sont devenus d'usage courant : réponse paralysable et non paralysable. Ces modèles représentent un comportement idéalisé, dont l'un ou l'autre ressemble souvent adéquatement à la réponse d'un système de comptage réel. Les hypothèses fondamentales des modèles sont représentées sur la figure 15. Au centre de cette figure, une échelle de temps est représentée sur laquelle sont indiqués six événements espacés de manière aléatoire dans le détecteur.

En bas de la figure (15) est représenté le comportement du temps mort correspondant d'un détecteur supposé non paralysable. Un temps fixe T est supposé suivre chaque vrai événement qui se

produit pendant la "période active" du détecteur. Les vrais événements qui se produisent pendant la période morte sont perdus et supposés n'avoir aucun effet sur le comportement du détecteur.

Dans l'exemple étudié, le détecteur non paralysable enregistrerait quatre comptages à partir des six vraies interactions. En revanche, le comportement d'un détecteur paralysable est représenté le long de l'axe supérieur en haut de la figure (15). Le même temps mort T est supposé suivre chaque véritable interaction qui se produit pendant la durée de fonctionnement du détecteur. Cependant, les événements réels qui se produisent pendant la période morte, bien qu'ils ne soient toujours pas enregistrés en tant que décomptes, sont supposés prolonger le temps mort d'une autre période T après l'événement perdu. Dans l'exemple de la figure (15), seuls trois comptages sont enregistrés pour les six événements réels.

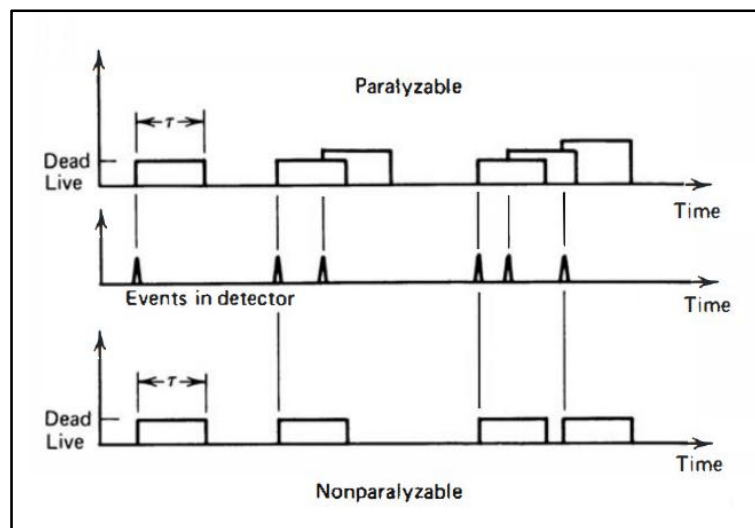


Fig. 15 Comportement de deux modèles du temps mort pour des détecteurs de rayonnement

Les deux modèles prédisent les mêmes pertes de premier ordre et ne diffèrent que lorsque les taux d'événements réels sont élevés. Ce sont en quelque sorte deux extrêmes du comportement idéalisé du système, et les systèmes de comptage réels afficheront souvent un comportement intermédiaire entre ces extrêmes. Le comportement détaillé d'un système de comptage spécifique peut dépendre des processus physiques se déroulant dans le détecteur lui-même ou des retards introduits par l'électronique de traitement et d'enregistrement des impulsions.

Par la suite, nous examinons la réponse d'un système de détection à une source de rayonnement en régime permanent, et nous adoptons les définitions suivantes:

n = taux d'interaction réel ;

m = taux de comptage enregistré ;

τ = temps mort du système.

Nous supposons que le temps de comptage est long, de sorte que n et m peuvent être considérés comme des taux moyens. En général, nous aimerions obtenir une expression du vrai taux d'interaction n en fonction du taux mesuré m et du temps mort du système τ , afin que les corrections appropriées puissent être apportées aux données mesurées pour tenir compte des pertes du temps mort.

Dans le cas non paralysable, la fraction de tout le temps pendant laquelle le détecteur est mort est donnée tout simplement par le produit $(m\tau)$. Par conséquent, la vitesse à laquelle les vrais événements sont perdus est simplement $(nm\tau)$. Puisque $(n - m)$ est une autre expression du taux de pertes,

$$\boxed{n - m = nm\tau} \quad (25)$$

En résolvant pour n , on obtient pour un modèle non paralysable :

$$\boxed{n = \frac{m}{1 - m\tau}} \quad (26)$$

Dans le cas paralysable, les périodes mortes ne sont pas toujours de durée fixe, on ne peut pas donc appliquer le même argument. Au lieu de cela, nous notons que le taux m est identique au taux d'occurrences d'intervalles de temps entre événements réels qui dépassent τ . La distribution des intervalles entre événements aléatoires se produisant à un taux moyen n :

$$\boxed{P_1(T) dT = ne^{-nT} dT} \quad (27)$$

Où $P_1(T) dT$ est la probabilité d'observer un intervalle dont la longueur est comprise entre dT autour de T .

La probabilité d'intervalles supérieurs à τ peut être obtenu en intégrant cette distribution entre τ et ∞ .

$$\boxed{P_2(\tau) = \int_{\tau}^{\infty} P_1(T) dT = e^{-n\tau}} \quad (28)$$

Le taux d'occurrence de tels intervalles est alors obtenu en multipliant simplement l'expression ci-dessus par le taux vrai n pour un modèle paralysable :

$$\boxed{m = ne^{-n\tau}} \quad (29)$$

Le modèle paralysable conduit à un résultat plus lourd car nous ne pouvons pas résoudre explicitement le vrai taux n . Au lieu de cela, l'équation précédente doit être résolue itérativement si n doit être calculé à partir des mesures de m et de la connaissance de τ .

Un tracé du taux observé m par rapport au taux réel (vrai) n est donné à la figure (15) pour les deux modèles. Lorsque les taux sont faibles, les deux modèles donnent pratiquement le même résultat, mais le comportement aux taux élevés est nettement différent. Un système non paralysable s'approchera d'une valeur asymptotique pour le taux observé de $1/\tau$, qui représente la situation dans laquelle le compteur a à peine le temps de terminer une période morte avant d'en entamer une autre. Pour le comportement paralysable, on voit que le taux observé passe par un maximum. Des taux d'interaction vrais très élevés entraînent une extension multiple de la période morte après un comptage initial enregistré, et très peu d'événements vrais peuvent être enregistrés. Il faut être toujours prudent lors de l'utilisation d'un système de comptage qui peut être paralysable pour s'assurer que des taux observés ostensiblement faibles correspondent en fait à des taux d'interaction faibles plutôt qu'à des taux très élevés du côté opposé au maximum. Des erreurs dans l'interprétation des données de comptage nucléaire des systèmes paralysables se sont produites dans le passé en négligeant le fait qu'il existe toujours deux taux d'interaction réels possibles correspondant à un taux observé donné. Comme le montre la figure (16), le taux observé m_1 peut correspondre aux taux réels n_1 ou n_2 . L'ambiguïté ne peut être résolue qu'en modifiant le taux réel dans une direction connue tout en observant si le taux observé augmente ou diminue.

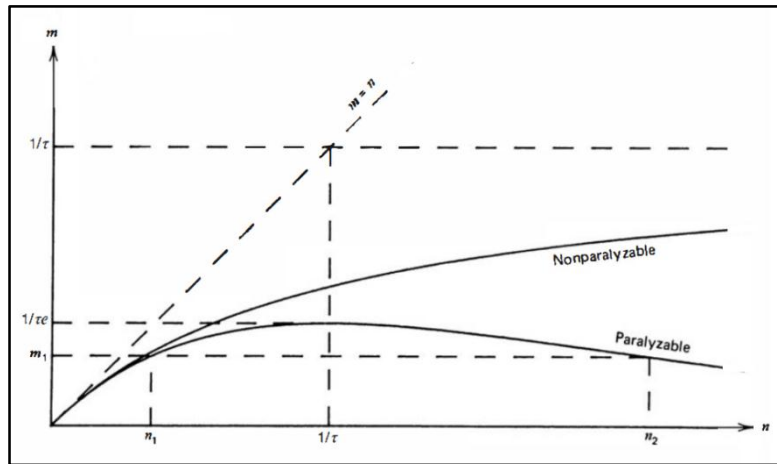


Fig. 16 Variation du taux observé m en fonction du taux réel n pour deux modèles de pertes du temps mort

Pour des débits taux ($n \ll 1/\tau$) les approximations suivantes peuvent s'écrire :

<i>Nonparalyzable</i>	$m = \frac{n}{1 + n\tau} \cong n(1 - n\tau)$	(30)
<i>Paralyzable</i>	$m = ne^{-n\tau} \cong n(1 - n\tau)$	

Ainsi, les deux modèles conduisent à des résultats identiques dans la limite des petites pertes du temps mort. Si c'est possible, il convient d'éviter les conditions de mesure dans lesquelles les pertes de temps mort sont élevées en raison des erreurs qui se produisent inévitablement lors de la correction des pertes. La valeur de τ peut-être incertaine ou sujette à variation, et le comportement du système peut ne pas suivre exactement l'un ou l'autre des modèles décrits ci-dessus. Lorsque les pertes sont supérieures à 30 ou 40 %, le taux réel calculé devient très sensible aux petites modifications du taux mesuré et du comportement supposé du système. Au lieu de cela, l'utilisateur doit chercher à réduire les pertes en changeant les conditions de mesure ou en choisissant un système de comptage avec un temps mort plus petit.

8.2 Méthodes de mesure du temps mort

Afin d'effectuer des corrections du temps mort à l'aide de l'un ou l'autre des modèles, une connaissance préalable du temps mort τ est nécessaire. Parfois, ce temps mort peut être associé à une propriété limite connue du système de comptage (par exemple, un temps de résolution fixe d'un circuit électronique). Le plus souvent, le temps mort ne sera pas connu ou pourra varier selon les conditions de fonctionnement et devra donc être mesuré directement. Les techniques de mesure courantes sont basées sur le fait que le taux observé varie de manière non linéaire avec le taux réel. Par conséquent, en supposant que l'un des modèles spécifiques est applicable, et en mesurant le taux observé pour au moins deux taux réels différents qui diffèrent par un rapport connu, le temps mort peut être calculé.

L'exemple courant est la méthode à deux sources. La méthode est basée sur l'observation du taux de comptage de deux sources individuellement et en combinaison. Comme les pertes de comptage ne sont pas linéaires, le taux observé dû aux sources combinées sera inférieur à la somme des taux dus aux deux sources comptées individuellement, et le temps mort peut être calculé à partir de l'écart. Pour illustrer la méthode, supposons que n_1 , n_2 et n_{12} soient les vrais taux de comptage (échantillon plus bruit de fond) avec la source (1), la source (2) et les deux sources combinées, respectivement. Soient m_1 , m_2 et m_{12} les taux observés correspondants. Aussi, soit n_b et m_b les taux vrais et de fond mesurés avec les deux sources supprimées. Alors :

$$\begin{aligned} n_{12} - n_b &= (n_1 - n_b) + (n_2 - n_b) \\ n_{12} + n_b &= n_1 + n_2 \end{aligned} \tag{31}$$

En supposant maintenant le modèle non paralysable et en substituant son équation, on obtient :

$$\frac{m_{12}}{1 - m_{12}\tau} + \frac{m_b}{1 - m_b\tau} = \frac{m_1}{1 - m_1\tau} + \frac{m_2}{1 - m_2\tau} \tag{32}$$

La résolution explicite de cette équation pour τ donne le résultat suivant :

$$\tau = \frac{X(1 - \sqrt{1 - Z})}{Y} \quad (33)$$

Avec :

$$X \equiv m_1 m_2 - m_b m_{12}$$

$$Y \equiv m_1 m_2 (m_{12} + m_b) - m_b m_{12} (m_1 + m_2)$$

$$Z \equiv \frac{Y(m_1 + m_2 - m_{12} - m_b)}{X^2}$$

Étant donné que la méthode est essentiellement basée sur l'observation de la différence entre deux grands nombres presque égaux, des mesures minutieuses sont nécessaires afin d'obtenir des valeurs fiables du temps mort. La mesure est généralement effectuée en comptant la source (1), en plaçant la source (2) à proximité et en mesurant le taux combiné, puis en supprimant la source 1 pour mesurer le taux causé par la source (2) seule. Au cours de cette opération, il faut veiller à ne pas déplacer la source déjà en place, et il faut tenir compte de la possibilité que la présence d'une deuxième source diffuse dans le détecteur un rayonnement qui ne serait normalement pas compté à partir de la première source seule. Afin de maintenir la diffusion inchangée, une deuxième source fictive sans activité est normalement mise en place lorsque les sources sont comptées individuellement. Les meilleurs résultats sont obtenus en utilisant des sources suffisamment actives pour aboutir à un temps mort $m_{12}\tau$ d'au moins 20 %.

Une deuxième méthode peut être appliquée si une source de radio-isotopes à courte durée de vie est disponible. Dans ce cas, l'écart entre le taux de comptage observé et la décroissance exponentielle connue de la source peut être utilisé pour calculer le temps mort. La technique, connue sous le nom de méthode de la source décroissante, est basée sur le comportement connu du taux vrai n :

$$n = n_0 e^{-\lambda t} + n_b \quad (34)$$

Où n_0 est le taux réel au début de la mesure et λ est la constante de désintégration de l'isotope particulier utilisé pour la mesure.

Dans la limite où le bruit de fond est négligeable, une procédure graphique simple peut être appliquée pour analyser les données résultantes. Puis l'équation précédente devient :

$$n \cong n_0 e^{-\lambda t} \quad (35)$$

En insérant la dernière équation dans l'équation représentant le modèle non paralysable, on obtient la relation suivante pour ce modèle :

$$me^{\lambda t} = -n_0\tau m + n_0 \quad (36)$$

Si nous définissons, sur la figure (17), l'abscisse comme m et l'ordonnée comme le produit $me^{\lambda t}$, alors la dernière équation est celle d'une ligne droite. La procédure expérimentale consiste à mesurer le taux observé m en fonction du temps t , et donc à définir des points qui doivent se situer sur cette droite en partant de la droite et en se déplaçant vers la gauche au fur et à mesure que la source se désintègre. En ajustant la meilleure ligne droite aux données, l'ordonnée à l'origine donnera n_0 , le taux vrai au début de la mesure, et la pente négative donnera le produit de $n_0\tau$. Le temps mort τ découle alors directement du rapport de la pente à l'interception.

Pour le modèle paralysable, en insérant l'équation $n \cong n_0e^{-\lambda t}$ dans l'équation représentant ce modèle on obtient le résultat suivant :

$$\lambda t + \ln m = -n_0\tau e^{-\lambda t} + \ln n_0 \quad (37)$$

Encore une fois en choisissant l'abscisse et l'ordonnée comme indiqué sur la figure (17 b), l'équation d'une ligne droite en résulte. Dans ce cas, l'ordonnée à l'origine donne la valeur $\ln n_0$, tandis que la pente donne à nouveau la valeur négative du produit $n_0\tau$. Le temps mort peut être simplement déduit de ces deux valeurs. La méthode de la source décroissante offre l'avantage non seulement de pouvoir mesurer la valeur du temps mort mais aussi de pouvoir tester la validité des modèles supposés. Si un modèle non paralysable décrit mieux le système de comptage, les données correspondront mieux à une ligne droite pour la forme illustrée à la figure (17 a). En revanche, si un modèle paralysable est plus approprié, la forme indiquée dans figure (17 b) donnera un tracé plus linéaire des données. Pour être efficaces, les mesures doivent être effectuées pendant une période de temps, au moins égale à la demi-vie du radio-isotope en décomposition, et le taux de perte initiale $m\tau$ doit être d'au moins 20 %.

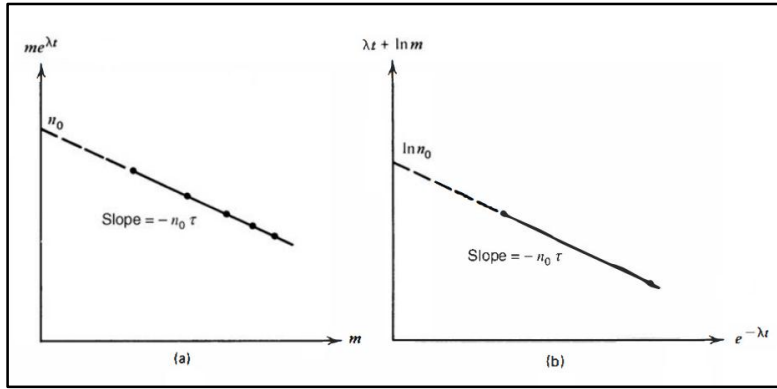


Fig. 17 Application de la méthode de la source décroissante pour déterminer le temps mort

Si le bruit de fond est supérieur à quelques pour cent du plus petit taux mesuré, la procédure graphique peut entraîner des erreurs importantes. Bien qu'une certaine amélioration résulte de la soustraction du taux de fond observé de toutes les valeurs de m mesurées, cette correction n'est pas rigoureuse et une analyse exacte ne peut être faite qu'en revenant à l'équation (34). Il devient alors nécessaire d'utiliser des techniques de calcul numérique pour obtenir des valeurs pour n_0 et τ qui, lorsqu'elles sont insérées dans un modèle supposé du comportement du temps mort, le résultat est un meilleur ajustement aux données mesurées.

Les méthodes décrites ci-dessus sont applicables pour déterminer le temps mort de systèmes de comptage simples dans lesquels toutes les impulsions supérieures à une amplitude seuil sont enregistrées. Pour les systèmes spectroscopiques dans lesquels le spectre de hauteur d'impulsion est enregistré à partir du détecteur, d'autres méthodes sont disponibles pour traiter les pertes de comptage basées sur le mélange d'impulsions artificielles d'un générateur d'impulsions dans la chaîne de traitement du signal.

8.3 Statistiques des pertes du temps mort

Lors de la mesure du rayonnement provenant de sources à l'état stable, nous supposons normalement que les véritables événements se produisant dans le détecteur suivent les statistiques de Poisson dans lesquelles la probabilité qu'un événement se produise par unité de temps est une constante. Le temps mort du système a pour effet de supprimer sélectivement certains des événements avant qu'ils ne soient enregistrés. Spécifiquement, les événements survenant après de courts intervalles de temps suivant des événements précédents sont préférentiellement rejetés, et la distribution d'intervalle, donnée par l'équation ci-dessous, est modifiée par rapport à sa forme exponentielle normale, figure (18). Les pertes de temps mort faussent donc les statistiques des comptages enregistrés loin d'un vrai comportement de Poisson. Si les pertes sont faibles ($n\tau$ inférieur à 10 ou 20 %), cette distorsion a cependant peu d'effet pratique sur la validité des formules statistiques de prédiction des incertitudes statistiques de comptage.

$$I_1(t) dt = re^{-rt} dt$$

(38)

$I_1(t)$ est la fonction de distribution des intervalles entre événements aléatoires adjacents.

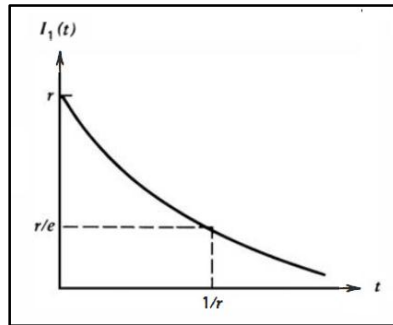


Fig. 18 Forme exponentielle de la distribution des intervalles entre événements aléatoires adjacents

Si les pertes de temps mort ne sont pas faibles, les écarts par rapport aux statistiques de Poisson deviennent plus significatifs. Le rejet d'événements qui se produisent après de courts intervalles de temps fait que la séquence des comptages enregistrés devient un peu plus régulière, et la variance attendue dans les mesures répétées est réduite.

Avec un comportement paralysable ou non paralysable, il y a une chance que plus d'un événement soit perdu par période morte. Un nombre enregistré peut donc correspondre à l'occurrence de n'importe quel nombre d'événements vrais. La probabilité relative que plusieurs événements vrais sont contenus dans une période morte augmentera à mesure que le taux d'événements vrais augmentera. Puisque les vrais événements obéissent toujours aux statistiques de Poisson, des analyses relativement simples peuvent être faites pour prédire la probabilité qu'un décompte enregistré résulte de la combinaison d'exactly x événements vrais.

8.4 Pertes de temps mort provenant de sources pulsées

L'analyse des pertes de comptage dues au temps mort du détecteur supposait que le détecteur était irradié par une source en régime permanent d'intensité constante. Il existe de nombreuses applications dans lesquelles la source de rayonnement n'est pas continue mais consiste plutôt en de courtes impulsions répétées à une fréquence constante. Par exemple, les accélérateurs linéaires d'électrons utilisés pour générer des rayons X à haute énergie peuvent être utilisés pour produire des impulsions d'une largeur de quelques microsecondes avec une fréquence de répétition de plusieurs kilohertz. Dans de tels cas, les résultats donnés précédemment peuvent ne pas être applicables pour corriger correctement les effets des pertes dues au temps mort. Il faut maintenant appliquer des méthodes analytiques de substitution qui utilisent certains des principes statistiques.

Nous limitons notre analyse à une source de rayonnement qui peut être représentée par l'intensité dépendante du temps donnée dans la figure (19) :

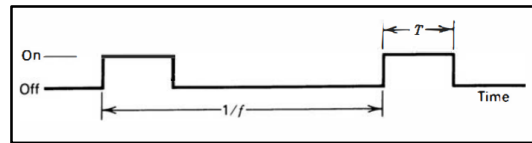


Fig. 19 Intensité de source de rayonnement émettant de courtes impulsions répétées à une fréquence constante

On suppose que l'intensité du rayonnement est constante pendant toute la durée T de chaque impulsion, et que les impulsions se produisent à une fréquence constante (f) Selon la valeur relative du temps mort du détecteur τ , plusieurs conditions peuvent s'appliquer :

1. Si τ est beaucoup plus petit que T , le fait que la source soit pulsée a peu d'effet, et les résultats donnés précédemment dans cette section pour les sources en régime permanent peuvent être appliqués avec une précision acceptable.
2. Si τ est inférieur à T mais pas avec une grande fraction, seul un petit nombre de comptages peut être enregistré par le détecteur au cours d'une seule impulsion. C'est la circonstance la plus compliquée et elle dépasse le cadre de la présente discussion.
3. Si τ est supérieur à T mais inférieur au temps "off" entre les impulsions (donné par « $1/f - T$ »), l'analyse suivante s'applique. A noter que dans ces conditions, on ne peut avoir au maximum qu'un seul comptage du détecteur par impulsion source. De plus, le détecteur sera entièrement récupéré au début de chaque impulsion.

Pour l'analyse qui va suivre, il n'est plus nécessaire que le rayonnement soit constant sur la durée d'impulsion T , mais il faut que chaque impulsion de rayonnement soit de même intensité. Nous suivons les définitions suivantes :

τ = temps mort du système de détection ;

m = taux de comptage observé ;

n = vrai taux de comptage (si τ était 0) ;

T = longueur d'impulsion source ;

f = fréquence d'impulsion de la source.

Puisqu'il ne peut y avoir qu'un seul comptage par impulsion, la probabilité d'un comptage observé par impulsion source est donnée par m/f .

Le nombre moyen d'événements vrais par impulsion source est par définition égal à n/f (notez que cette moyenne peut être supérieure à 1.) On peut appliquer la distribution de Poisson :

$$P(x) = \frac{(\bar{x})^x e^{-\bar{x}}}{x!} \quad (39)$$

Pour prédire la probabilité qu'au moins un événement vrai se produit par impulsion source :

$$P(> 0) = 1 - P(0) = 1 - e^{-\bar{x}} = 1 - e^{-n/f} \quad (40)$$

Étant donné que le détecteur est "actif" au début de chaque impulsion, un comptage sera enregistré si au moins un événement réel se produit pendant l'impulsion. Un seul comptage de ce type peut être enregistré, de sorte que l'expression ci-dessus est également la probabilité d'enregistrer un comptage par impulsion source. En égalant les deux expressions de cette probabilité, on obtient :

$$\frac{m}{f} = 1 - e^{-n/f} \quad (41)$$

Ou bien

$$m = f(1 - e^{-n/f}) \quad (42)$$

La courbe qui caractérise ce comportement (variation de m en fonction de n) est donnée par la figure (20) :

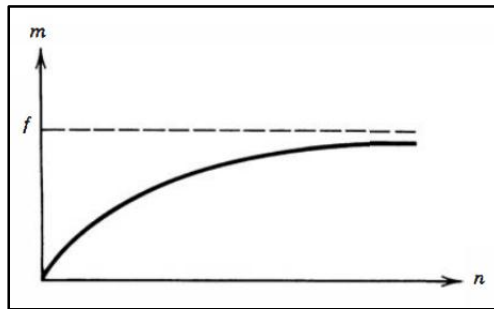


Fig. 20 Probabilité d'enregistrer un comptage par impulsion source

Dans ces conditions, le taux de comptage maximal observable est simplement la fréquence de répétition des impulsions, puisque pas plus d'un seul comptage ne peut être enregistré par impulsion. En outre, ni la durée spécifique du temps mort ni le comportement détaillé du temps mort du système (qu'il soit paralysable ou non paralysable) n'a aucune influence sur les pertes.

Dans des circonstances normales, nous sommes plus intéressés par une formule de correction pour prédire le débit réel à partir du débit mesuré et du temps mort du système. La résolution de l'équation de m pour n , conduit à :

$$n = f \ln \left(\frac{f}{f - m} \right) \quad (43)$$

Rappelons que cette correction n'est valable que sous les conditions $T < \tau < (1/f - T)$.

Dans ce cas, les pertes dues au temps mort sont faibles sous la condition $m \ll f$. En développant le terme logarithmique ci-dessus pour cette limite, on trouve qu'une correction du premier ordre est alors donnée par :

$$n \cong \frac{m}{1 - m/2f} \quad (44)$$

Ce résultat, en raison de sa similarité avec l'équation (26) peut être considéré comme prédisant d'une valeur de temps mort effectif de $1/2f$ dans cette limite de perte faible. Étant donné que cette valeur est maintenant la moitié de la période d'impulsion de la source, elle peut être plusieurs fois supérieure au temps mort physique réel du détecteur.

Introduction à la détection et les détecteurs de rayonnements

1. Principe de la détection

Le principe de base de la détection consiste à faire interagir la particule avec la matière pour lui prélever l'ensemble ou une partie de son énergie initiale. Celle-ci est la plupart du temps transformée en un signal électrique qui va être ensuite utilisé pour obtenir toutes les informations sur la particule.

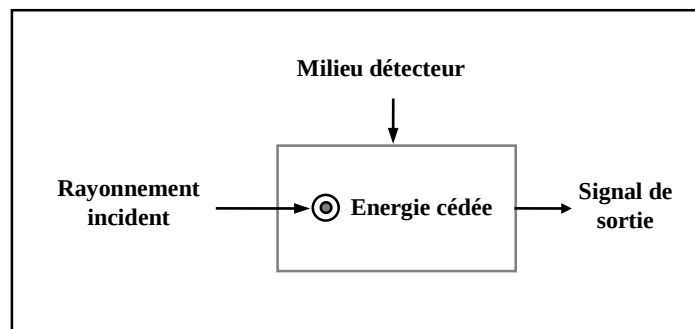


Fig. 01 : Principe de la détection de rayonnement

2. Détecteur de rayonnements

Un détecteur va donc être avant tout constitué de matière qui sera mis sur le parcours de la particule. Cette matière pourra être solide, liquide ou gaz. Le choix et la nature de la matière ne dépendront pas que des différents rayonnements à étudier mais aussi de leurs grandeurs physiques à mesurer. Le processus de transformation de l'énergie cédée par la particule en un signal électrique impose aussi le choix des matériaux utilisés pour construire les détecteurs.

Trois grands types de détecteurs sont à considérer :

- Les détecteurs à gaz ;
 - Chambre d'ionisation ;
 - Compteurs proportionnels ;
 - Tubes Geiger.
- Les scintillateurs ;
- Les détecteurs à base de semi-conducteurs.

Annexe I

Détecteurs à gaz (Chambre d'ionisation)

1. Introduction

Plusieurs des types de détecteurs de rayonnement sont basés sur les effets produits lorsqu'une particule chargée traverse un gaz. Les principaux modes d'interaction impliquent l'ionisation et l'excitation des molécules de gaz le long de la trajectoire des particules. Bien que les molécules excitées puissent parfois être utilisées pour générer un signal électrique, la majorité des détecteurs à gaz sont basés sur l'ionisation créée par le passage du rayonnement pour produire un signal électrique. Les détecteurs qui font l'objet des trois parties de cet annexe (chambres d'ionisation, compteurs proportionnels, tubes Geiger) délivrent tous, de manière quelque peu différente, un signal de sortie électrique qui provient des paires d'ions formées dans le gaz remplissant le détecteur.

Les chambres d'ionisation sont en principe les plus simples de tous les détecteurs à gaz. Leur fonctionnement normal repose sur la collecte de toutes les charges créées par ionisation directe au sein du gaz grâce à l'application d'un champ électrique. Comme pour les autres détecteurs, les chambres d'ionisation peuvent fonctionner en mode courant ou en mode impulsion. Dans la plupart des applications courantes, les chambres d'ionisation sont utilisées en mode courant, bien que dans quelques d'applications ils sont utilisés en mode impulsif. En revanche, les compteurs proportionnels ou les tubes Geiger sont presque toujours utilisés en mode impulsif.

Le terme chambre d'ionisation est utilisé exclusivement pour le type de détecteur dans lequel des paires d'ions sont collectées à partir de gaz. Le processus correspondant dans les solides est la collection de paires électron-trou dans les détecteurs à semi-conducteurs.

2. Processus d'ionisation dans les gaz

Lorsqu'une particule chargée rapide traverse un gaz, ses interactions créent à la fois des molécules excitées et des molécules ionisées le long de sa trajectoire. Après qu'une molécule neutre est ionisée, l'ion positif résultant et l'électron libre sont appelés une paire d'ions, qui sert de constituant de base du signal électrique développé par la chambre d'ionisation. Les ions peuvent être formés soit par interaction directe avec la particule incidente, soit par un processus secondaire dans lequel une partie de l'énergie de la particule est d'abord transférée à un électron énergétique ou "rayonnement delta". Indépendamment des mécanismes possibles, la quantité pratique essentielle est le nombre total de paires d'ions créées le long de la trajectoire du rayonnement.

Tab. 01 : Valeurs de l'énergie dissipée par paire d'ions pour différents gaz

Values of the Energy Dissipation per Ion Pair (the W-Value) for Different Gases ^a			
Gas	First Ionization Potential (eV)	W-Value (eV/ion pair)	
		Fast Electrons	Alpha Particles
Ar	15.7	26.4	26.3
He	24.5	41.3	42.7
H ₂	15.6	36.5	36.4
N ₂	15.5	34.8	36.4
Air		33.8	35.1
O ₂	12.5	30.8	32.2
CH ₄	14.5	27.3	29.1

^aValues for W from ICRU , "Average Energy Required to Produce an Ion Pair," International Commission on Radiation Units and Measurements, Washington, DC, 1979.

2.1 Nombre de paires d'ions formées

Au minimum, la particule doit transférer une quantité d'énergie égale à l'énergie d'ionisation de la molécule de gaz pour permettre au processus d'ionisation de se produire. Dans la plupart des gaz, l'énergie d'ionisation pour les couches d'électrons les moins liées est comprise entre 10 et 25 eV. Cependant, il existe d'autres mécanismes par lesquels la particule incidente peut perdre de l'énergie dans le gaz qui ne crée pas d'ions. Ce sont des processus d'excitation dans lesquels un électron peut être élevé à un état lié supérieur dans la molécule sans être complètement éliminé. Par conséquent, l'énergie moyenne perdue par la particule incidente par paire d'ions formée (définie comme la valeur W) est toujours sensiblement supérieure à l'énergie d'ionisation. La valeur W est en principe fonction de l'espèce de gaz impliquée, du type de rayonnement et de son énergie. Cependant, les observations empiriques, montrent qu'il ne s'agit pas d'une fonction forte d'aucune de ces variables et qu'il s'agit d'un paramètre constant pour de nombreux gaz et différents types de rayonnement. (Cette propriété des gaz contraste avec le comportement observé pour les milieux détecteurs à scintillation et à semi-conducteurs discutés plus loin, qui présentent tous deux de grandes différences dans la quantité de signal généré par différents types de particules de même énergie.) Certaines données spécifiques sont présentées dans Tableau 01, et une valeur typique est 25-35 eV/paire d'ions. Par conséquent, une particule incidente de 1 MeV, si elle est complètement arrêtée dans le gaz, créera environ 30 000 paires d'ions. En supposant que W est constante pour un type de rayonnement donné, l'énergie déposée sera proportionnelle au nombre de paires d'ions formées et pourra être déterminée si une mesure correspondante du nombre de paires d'ions est effectuée.

2.2 Facteur de Fano

Outre le nombre moyen de paires d'ions formées par chaque particule incidente, la fluctuation du leur nombre pour des particules incidentes d'énergie identique est également intéressante. Ces fluctuations fixeront une limite fondamentale à la résolution en énergie qui peut être obtenue dans n'importe quel détecteur basé sur la collecte des ions. Dans le modèle le plus simple, la formation de chaque paire d'ions sera considérée comme un processus de Poisson. Le nombre total de paires d'ions formées doit donc être soumis à des fluctuations caractérisées par un écart-type égal à la racine carrée du nombre moyen formé. De nombreux détecteurs de rayonnement présentent une fluctuation inhérente inférieure à celle prédite par ce modèle simplifié. Le facteur de Fano est introduit comme une constante empirique par laquelle la variance prédite doit être multipliée pour donner la variance observée expérimentalement.

Le facteur Fano reflète dans une certaine mesure la fraction de l'énergie des particules incidentes qui est convertie en charges (supports d'informations) dans le détecteur. Si toute l'énergie du rayonnement incident est convertie en paires d'ions et que la même quantité d'énergie est nécessaire pour former chaque paire d'ions, le nombre de paires produites serait toujours exactement le même et il n'y aurait pas de fluctuation statistique. Dans ces conditions, le facteur Fano serait nul.

L'introduction du facteur Fano souligne également l'importance des statistiques de partitionnement lorsqu'une partie de l'énergie des particules est perdue au profit de processus qui ne produit pas de paires d'ions. Cependant, si seule une très petite fraction du rayonnement incident est convertie, les paires d'ions se formeraient loin les unes des autres et avec une probabilité relativement faible, et il y aurait une bonne raison de s'attendre à ce que la distribution de leur nombre suive une loi de distribution de Poisson. Dans les gaz, on observe empiriquement que le facteur Fano est inférieur à 1, de sorte que les fluctuations sont plus petites que ce qui serait prédit sur la base des seules statistiques de Poisson.

Le facteur Fano n'a de signification que lorsque le détecteur fonctionne en mode impulsif. Nous reportons donc les discussions sur son ampleur dans les gaz au chapitre suivant sur les proportionnels, où le fonctionnement en mode impulsif et une bonne résolution en énergie sont des considérations plus importantes.

2.3 Diffusion, transfert de charge et recombinaison

Les atomes ou molécules neutres du gaz sont en mouvement thermique constant, caractérisé par un libre parcours moyen pour les gaz typiques dans des conditions standard d'environ 10^{-6} à 10^{-8} m. Les ions positifs ou les électrons libres créés dans le gaz participent également au mouvement thermique aléatoire et ont donc une certaine tendance à diffuser loin des régions de haute densité. Ce processus de diffusion est beaucoup plus important pour les électrons libres que pour les ions puisque

leur vitesse thermique moyenne est beaucoup plus grande. Une collection ponctuelle d'électrons libres se propagera autour du point d'origine dans une distribution spatiale gaussienne dont la largeur augmentera avec le temps. Si nous laissons σ soit l'écart type de cette distribution telle que projetée sur un axe orthogonal arbitraire (x, y ou z), et t le temps écoulé, alors on peut montrer que :

$$\sigma = \sqrt{2Dt} \quad (01)$$

La valeur du coefficient de diffusion D dans des cas simples peut être prédite à partir de la théorie cinétique des gaz, mais en général, un modèle de transport plus complexe est nécessaire pour modéliser avec précision les observations expérimentales.

Parmi les nombreux types de collisions qui se produisent entre les électrons libres, les ions et les molécules de gaz neutre, plusieurs qui sont importants pour comprendre le comportement des détecteurs remplis de gaz sont illustrés schématiquement dans la figure (01). Des collisions de transfert de charge peuvent se produire lorsqu'un ion positif rencontre une autre molécule de gaz neutre. Dans une telle collision, un électron est transféré de la molécule neutre à l'ion, inversant ainsi les rôles de chacun. Ce transfert de charge est particulièrement important dans les mélanges gazeux contenant plusieurs espèces moléculaires différentes. Il y aura alors une tendance à transférer la charge positive nette au gaz avec une énergie d'ionisation plus faible parce que l'énergie est libérée dans les collisions qui laissent cette espèce comme ion positif.

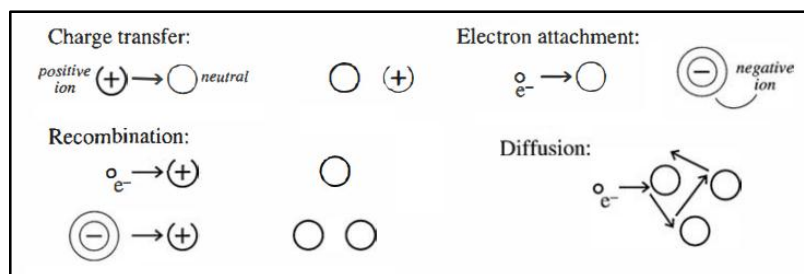


Fig. 01 Types d'interactions d'espèces chargées dans un gaz (les espèces en interaction sont à gauche et les produits de l'interaction sont à droite, les atomes ou molécules neutres sont représentés par de simples cercles)

L'électron libre membre de la paire d'ions d'origine subit également de nombreuses collisions dans sa diffusion normale. Dans certaines espèces de gaz, il peut y avoir une tendance à former des ions négatifs par la fixation de l'électron libre à une molécule de gaz neutre. Cet ion négatif partage alors de nombreuses propriétés avec l'ion positif d'origine formé dans le processus d'ionisation, mais avec une charge électrique opposée. L'oxygène est un exemple de gaz qui fixe facilement les électrons, de sorte que les électrons libres diffusant dans l'air sont rapidement convertis en ions négatifs. En revanche, l'azote, l'hydrogène, les gaz d'hydrocarbures et les gaz nobles sont tous

caractérisés par des coefficients d'attachement d'électrons relativement faibles, et donc l'électron continue de migrer dans ces gaz en tant qu'électron libre sous des conditions normales.

Les collisions entre les ions positifs et les électrons libres peuvent entraîner une recombinaison dans laquelle l'électron est capturé par l'ion positif et le ramène à un état de neutralité de charge. Alternativement, l'ion positif peut subir une collision avec un ion négatif dans laquelle l'électron supplémentaire est transféré à l'ion positif et les deux ions sont neutralisés. Dans les deux cas, la charge représentée par la paire d'origine est perdue et ne peut plus contribuer au signal dans les détecteurs basés sur la collecte de la charge d'ionisation.

Comme la fréquence de collision est proportionnelle au produit des concentrations des deux espèces impliquées, le taux de recombinaison peut s'écrire :

$$\boxed{\frac{dn^+}{dt} = \frac{dn^-}{dt} = -\alpha n^+ n^-} \quad (02)$$

Où :

n^+ = densité numérique des espèces positives ;

n^- = densité numérique des espèces négatives ;

α = coefficient de recombinaison.

Le coefficient de recombinaison est normalement de plusieurs ordres de grandeur plus grand entre les ions positifs et les ions négatifs par rapport à celui entre les ions positifs et les électrons libres. Dans les gaz qui forment facilement des ions négatifs par fixation d'électrons, pratiquement toute la recombinaison a lieu entre les ions positifs et négatifs.

Il existe deux types généraux de perte de recombinaison : la recombinaison colonnaire et la recombinaison volumique. Le premier type (parfois aussi appelé recombinaison initiale) provient du fait que des paires d'ions sont d'abord formées dans une colonne le long de la trajectoire de la particule ionisante. La densité locale de paires d'ions est donc élevée le long de la trajectoire jusqu'à ce que les paires d'ions soient amenées à dériver ou à diffuser hors de leur point de formation. La recombinaison colonnaire est la plus sévère pour les particules densément ionisantes telles que les particules alpha ou les ions lourds par rapport aux électrons rapides qui déposent leur énergie sur une trajectoire beaucoup plus longue. Ce mécanisme de perte dépend uniquement des conditions locales le long des pistes individuelles et ne dépend pas de la vitesse à laquelle ces pistes se forment dans le volume du détecteur. En revanche, le volume la recombinaison est due aux rencontres entre ions et / ou électrons après qu'ils aient quitté l'emplacement immédiat de la piste. Étant donné que de nombreuses traces se forment généralement au cours du temps nécessaire aux ions pour dériver jusqu'aux électrodes collectrices, il est possible que des ions et / ou les électrons de pistes indépendantes entrent en

collision et se recombinent. La recombinaison volumique augmente donc en importance avec le taux d'irradiation. Ainsi, la séparation et la collecte des charges doivent être aussi rapides que possible afin de minimiser la recombinaison et des champs électriques d'intensité élevée sont suggérés à cet égard.

2.4 Migration et collecte de charges

2.4.1 Mobilité des charges

Si un champ électrique externe est appliqué à la région dans laquelle des ions ou des électrons existent dans le gaz, les forces électrostatiques auront tendance à éloigner les charges de leur point d'origine. Le mouvement net consiste en une superposition d'une vitesse thermique aléatoire avec une vitesse de dérive nette dans une direction donnée. La vitesse de dérive des ions positifs est dans la direction du champ électrique conventionnel, tandis que les électrons libres et les ions négatifs dérivent dans la direction opposée.

Pour les ions dans un gaz, la vitesse de dérive peut être prédite assez précisément à partir de la relation :

$$v = \frac{\mu \mathcal{E}}{P} \quad (03)$$

Où :

v = vitesse de dérive

$\mathcal{E}$ = intensité du champ électrique

P = pression du gaz

La mobilité μ a tendance à rester assez constante sur de larges plages du champ électrique et de la pression du gaz et ne diffère pas beaucoup pour les ions positifs ou négatifs dans le même gaz. Les valeurs de la mobilité des ions sont comprises entre 1 et $1,5 \times 10^{-4} \text{ m}^2 \text{ atm} / \text{V} \cdot \text{s}$ pour les gaz de numéro atomique moyen. Par conséquent, à une pression de 1 atm, un champ électrique typique de 10^4 V/m entraînera une vitesse de dérive de l'ordre de 1 m/s. Les temps de transit des ions sur des dimensions d'un détecteur d'un centimètre seront donc d'environ 10 ms. Selon la plupart des normes, c'est une très longue période.

Les électrons libres se comportent tout à fait différemment, leur masse beaucoup plus faible leur permet une plus grande accélération entre les rencontres avec les molécules neutres du gaz, et la valeur de la mobilité dans l'équation (03) est typiquement 1000 fois supérieure à celle des ions. Les temps de collecte typiques pour les électrons sont donc de l'ordre des microsecondes, plutôt que des millisecondes. Dans certains gaz (les hydrocarbures et les mélanges argon-hydrocarbures) il existe un effet de saturation de la vitesse de dérive électrique. Sa valeur tend vers un maximum pour des valeurs élevées du champ électrique et peut même diminuer légèrement si le champ augmente encore.

Dans de nombreux autres gaz, la vitesse de dérive des électrons continue d'augmenter pour les plus grandes valeurs de $\mathcal{E}/p$ susceptibles d'être utilisées dans les compteurs remplis de gaz.

Lorsque les électrons dérivent à travers le gaz sous l'influence du champ électrique, ils suivront en première approximation le chemin de la ligne de champ électrique qui passe par leur point d'origine (mais dans le sens inverse du vecteur champ électrique).

Cependant, la diffusion aléatoire des électrons aura toujours lieu, ce qui obligera chaque électron individuel à suivre un chemin légèrement différent. Pour de fortes valeurs du champ électrique, l'augmentation de l'énergie moyenne donnée aux électrons dans la direction du champ donne différentes valeurs du coefficient de diffusion D dans l'équation (01) pour les directions parallèles ou transversales au champ. Au cours des quelques microsecondes qui sont généralement nécessaires pour que les électrons atteignent une électrode collectrice, la diffusion dans l'une ou l'autre direction peut être de l'ordre du millimètre ou moins. Alors que dans la plupart des détecteurs standard, cet étalement de charge aura peu d'effet pratique, il joue un rôle potentiel en limitant la résolution de position pouvant être atteinte dans les détecteurs à gaz "sensibles à la position" qui déduisent l'emplacement de l'événement ionisant à travers la position d'arrivée des électrons à l'anode, ou en mesurant le temps de dérive des électrons.

2.4.2 Le courant d'ionisation

En présence d'un champ électrique, la dérive des charges positives et négatives représentées par les ions et les électrons constitue un courant électrique. Si un volume donné de gaz subit une irradiation en régime permanent, la vitesse de formation des paires d'ions est constante. Pour tout petit volume d'essai du gaz, cette vitesse de formation sera exactement équilibrée par la vitesse à laquelle les paires d'ions sont perdues du volume, soit par recombinaison, soit par diffusion ou migration à partir du volume. Dans les conditions où la recombinaison est négligeable et toutes les charges sont efficacement collectées, le courant en régime permanent produit est une mesure précise de la vitesse à laquelle les paires d'ions se forment dans le volume. La mesure de ce courant d'ionisation est le principe de base de la chambre d'ionisation.

La figure (02) illustre les éléments de base d'une chambre d'ionisation rudimentaire. Un volume de gaz est enfermé dans une région dans laquelle un champ électrique peut être créé par l'application d'une tension externe. A l'équilibre, le courant circulant dans le circuit extérieur sera égal au courant d'ionisation collecté aux électrodes, et un ampèremètre sensible placé dans le circuit extérieur pourra donc mesurer le courant d'ionisation.

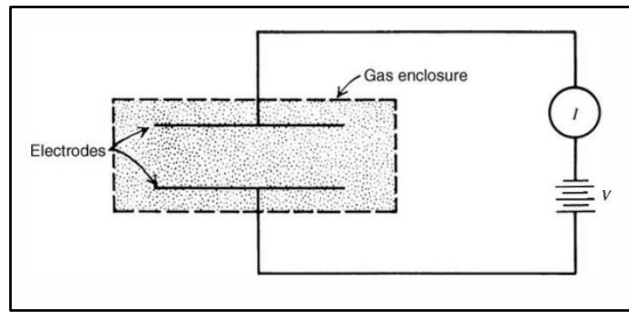


Fig. 02 Eléments de base d'une chambre d'ionisation

Les caractéristiques courant-tension d'une telle chambre sont également esquissées à la figure (03). En négligeant certains effets subtils liés aux différences des caractéristiques de diffusion entre les ions et les électrons, aucun courant net ne devrait circuler en l'absence d'une tension appliquée car aucun champ électrique n'existera alors dans le gaz. Les ions et électrons créés finissent par disparaître soit par recombinaison soit par diffusion du volume actif. Au fur et à mesure que la tension augmente, le champ électrique résultant commence à séparer les paires d'ions plus rapidement, et la recombinaison colonnaire diminue. Les charges positives et négatives sont également balayées vers les électrodes respectives avec une vitesse de dérive croissante, réduisant la concentration d'équilibre des ions dans le gaz et supprimant ainsi davantage la recombinaison volumique entre le point d'origine et les électrodes collectrices. Le courant mesuré augmente donc avec la tension appliquée car ces effets réduisent la quantité de charge d'origine qui est perdue. A un niveau suffisamment élevé de la tension appliquée, le champ électrique est suffisamment grand pour réduire efficacement la recombinaison à un niveau négligeable, et toutes les charges d'origine créées par le processus d'ionisation contribuent au courant ionique. Augmenter davantage la tension ne peut pas augmenter le courant car toutes les charges sont déjà collectées et leur taux de formation est constant. Il s'agit de la région de saturation ionique dans laquelle les chambres à ions sont classiquement exploitées. Dans ces conditions, le courant mesuré dans le circuit externe est une véritable indication du taux de formation de toutes les charges dues à l'ionisation dans le volume actif de la chambre.

À condition que la recombinaison soit négligeable, il convient de noter que le courant saturé ne change pas si les électrons d'ionisation restent sous forme d'électrons libres pendant leur dérive, ou sont au contraire attachés à des molécules de gaz pour former des ions négatifs. Si l'attachement se produit, la vitesse de dérive est maintenant des milliers de fois plus lente, mais la concentration à l'équilibre des charges négatives sera alors plus élevée du même facteur. Puisque le courant dépend du produit de la densité de charge et de la vitesse de dérive, le courant reste le même que si des électrons libres dériveraient.

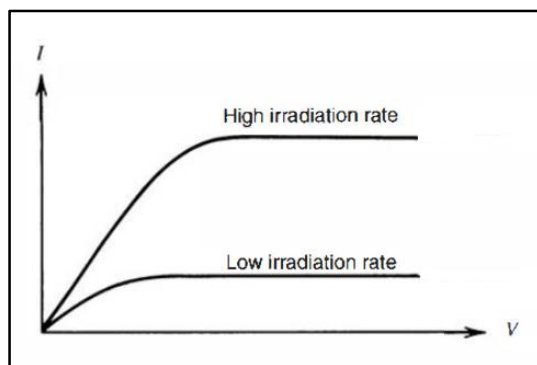


Fig. 03 Caractéristiques courant-tension d'une chambre d'ionisation

2.4.3 Facteurs affectant le courant de saturation

Plusieurs facteurs peuvent empêcher d'atteindre le courant de saturation dans une chambre d'ionisation. Le plus important d'entre eux est la recombinaison, qui est minimisée en s'assurant qu'une grande valeur du champ électrique existe partout dans le volume de la chambre d'ionisation. La recombinaison colonnaire le long de la trajectoire des particules chargées lourdes (telles que les particules alpha ou les fragments de fission) est particulièrement importante, de sorte que des valeurs plus importantes de la tension appliquée sont nécessaires dans ces cas pour atteindre la saturation, par rapport à l'irradiation par électrons ou rayons gamma. De plus, les effets de la recombinaison volumique sont plus importants pour une intensité d'irradiation plus élevée (et donc des courants ioniques plus élevés), comme illustré à la figure (04).

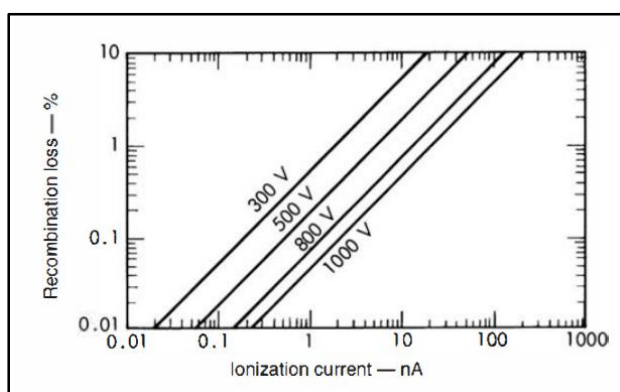


Fig. 04 Pertes dues à la recombinaison dans une chambre ionique remplie d'argon à 1 atm. Ces pertes sont minimisées à des valeurs élevées de la tension comme indiqué

2.5 Conception et fonctionnement des chambres d'ionisation DC

2.5.1 Considérations générales

Lorsqu'une chambre ionique est utilisée en mode courant continu (DC), il est possible de collecter les charges négatives soit sous forme d'électrons libres, soit sous forme d'ions négatifs. Par conséquent, pratiquement n'importe quel gaz peut être utilisé pour la chambre, y compris ceux qui ont des coefficients d'attachement d'électrons élevés. Bien que la recombinaison soit plus importante lors de la formation d'ions négatifs, les pertes par diffusion sont moins importantes. Les conditions de saturation peuvent être normalement atteintes sur des dimensions de quelques centimètres en utilisant des tensions appliquées qui ne dépassent pas des dizaines ou des centaines de volts. L'air est le gaz de remplissage le plus courant dans lequel les ions négatifs se forment facilement. L'air est nécessaire dans les chambres conçues pour la mesure de l'exposition aux rayons gamma. Des gaz plus denses tels que l'argon sont parfois choisis dans d'autres applications pour augmenter la densité d'ionisation dans un volume donné. La pression du gaz de remplissage est souvent de 1 atmosphère, bien que des pressions plus élevées soient parfois utilisées pour augmenter la sensibilité.

La géométrie choisie pour une chambre ionique peut varier considérablement en fonction de l'application, à condition que le champ électrique dans tout le volume actif puisse être maintenu suffisamment élevé pour conduire à la saturation du courant ionique. Des plaques parallèles ou une géométrie plane conduit à un champ électrique uniforme entre les plaques. Une géométrie cylindrique est également courante dans laquelle l'enveloppe extérieure du cylindre fonctionne au potentiel de terre et une tige conductrice axiale transporte la tension appliquée. Dans ce cas, un champ variant inversement au rayon est créé. Plusieurs méthodes analytiques peuvent être utilisées pour prédire les caractéristiques courant-tension des chambres de différentes géométries.

2.5.2 Isolateurs et anneaux de garde

Quelle que soit la conception, un isolant de support doit être prévu entre les deux électrodes. Étant donné que les courants d'ionisation typiques dans la plupart des applications sont extrêmement faibles (de l'ordre de 10^{-12} A ou même moins), le courant de fuite à travers ces isolateurs doit être maintenu très faible. Dans le schéma simple de la figure (02), toute fuite à travers l'isolant s'ajoutera au courant d'ionisation mesuré comme une composante indésirable du signal. Pour maintenir cette composante en dessous de 1 % du courant d'ionisation de 10^{-12} A pour une tension appliquée de 100 V, il faudrait que la résistance de l'isolant soit supérieure à 10^{16} ohms.

Un schéma illustrant l'utilisation d'un anneau de garde est présenté à la figure (05). L'isolant est maintenant segmenté en deux parties, une partie séparant l'anneau de garde conducteur de l'électrode négative et l'autre partie le séparant de l'électrode positive.

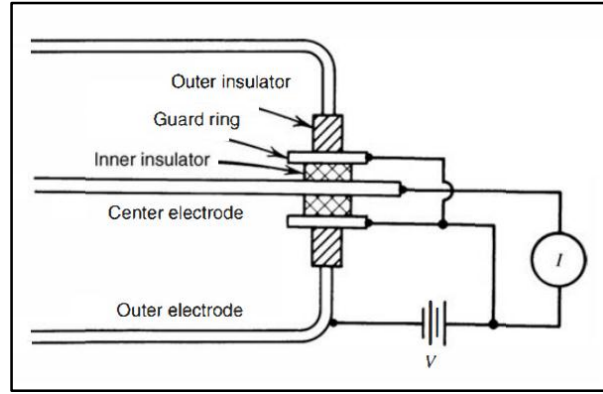


Fig. 05 Coupe transversale d'une extrémité d'une chambre d'ionisation cylindrique utilisant une construction à anneau de garde. La majeure partie de la tension appliquée V apparaît aux bornes de l'isolant extérieur, pour lequel le courant de fuite résultant ne contribue pas au courant mesuré I .

2.5.3 Mesure du courant ionique

L'amplitude du courant d'ionisation dans des conditions typiques est trop faible pour être mesurée à l'aide de techniques de galvanomètre standard. Au lieu de cela, une amplification active du courant doit être effectuée pour permettre sa mesure indirecte. Un électromètre mesure indirectement le courant en détectant la chute de tension aux bornes d'une résistance série placée dans le circuit de mesure (voir figure 06). La tension développée aux bornes de la résistance (typiquement avec une valeur de $10^9 - 10^{12}$ ohms) peut être amplifiée et sert de base au signal mesuré. Un point faible des circuits d'électromètre standard est qu'ils doivent être entièrement découplés. Toute petite dérive ou modification progressive des valeurs des composants entraîne donc une modification correspondante du courant de sortie mesuré. Ainsi, les circuits de ce type doivent être fréquemment équilibrés en court-circuitant l'entrée et en remettant l'échelle à zéro.

Une approche alternative consiste, au début, à convertir le signal continu (DC) en signal alternatif (AC), ce qui permet ensuite une amplification plus stable du signal alternatif. Cette conversion est accomplie dans une capacité dynamique ou dans un électromètre à ancre vibrante en collectant courant ionique à travers un circuit RC avec une large constante de temps, comme illustré à la figure (07). A l'équilibre, une tension constante est développée à travers ce circuit et qui est donnée par :

$$V = IR$$

(04)

Où I est le courant d'ionisation en régime permanent. Une charge Q est stockée dans la capacité et qui est donnée par :

$$Q = CV$$

(05)

Si la capacité est maintenant amenée à se changer rapidement par rapport à la constante de temps du circuit, un changement correspondant sera induit dans la tension aux bornes de C et qui est donnée par :

$$\Delta V = \frac{Q}{C^2} \Delta C \quad (06)$$

$$\Delta V = I \frac{R}{C} \Delta C \quad (07)$$

Si l'on fait varier sinusoïdalement la valeur de la capacité autour d'une valeur moyenne C, l'amplitude de la tension alternative induite est donc proportionnelle au courant d'ionisation.

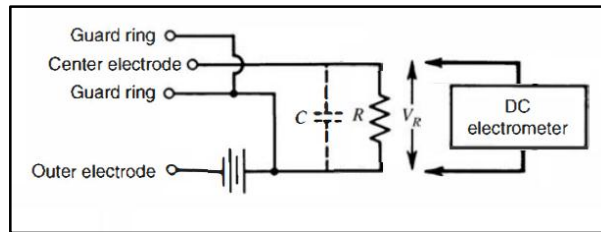


Fig. 06 Mesure de faibles courants ioniques grâce à l'utilisation d'une résistance série R et d'un électromètre pour enregistrer la tension résultante V_R . La capacité de la chambre avec toute capacité parasite parallèle est représentée par C. À condition que le courant ionique ne change pas pour plusieurs valeurs de la mesure constante de temps RC, sa valeur en régime permanent est donnée par : $I = V_R/R$

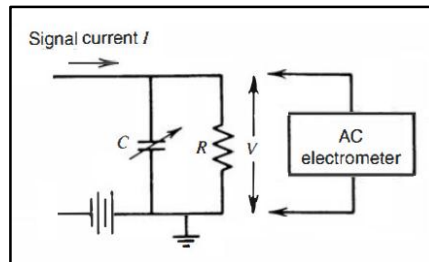


Fig. 07 Principe de l'électromètre à anse vibrante. Les oscillations de la capacité induisent une tension alternative proportionnelle au signal du courant en régime permanent I

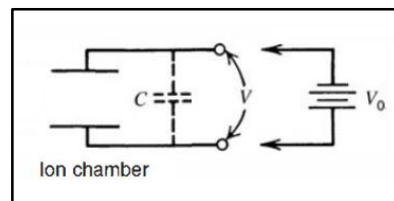


Fig. 08 Mesure de la charge totale due à l'ionisation ΔQ

La mesure de la charge totale due à l'ionisation ΔQ est effectuée sur une durée finie par intégration du courant. La tension aux bornes de la chambre est initialement fixée à V_0 et la source de tension est déconnectée. Ensuite en mesurant la tension V , la charge d'ionisation est donnée par :

$$\Delta Q = C \Delta V \quad (08)$$

Où

$$\Delta V = V_0 - V \quad (09)$$

Et C , est la capacité de la chambre avec toute capacité en parallèle.

Le courant d'ionisation moyen peut être également mesuré sur des périodes de temps finies par des méthodes d'intégration. Si la valeur de R sur la figure (06) est rendue infinie, tout courant d'ionisation de la chambre est simplement intégré à travers la capacité C . En notant le changement de la tension aux bornes de la capacité pendant la période de mesure, on peut en déduire le courant d'ionisation intégré total ou la charge d'ionisation. Si la perte à travers la capacité peut être maintenue faible, cette technique d'intégration a le potentiel de pouvoir mesurer des courants d'ionisation beaucoup plus faibles que par une mesure directe du courant.

Nous supposons que la chambre d'ionisation de la figure (08) est initialement chargée à une tension V_0 . Si les pertes à travers les isolateurs de la chambre et le condensateur externe sont négligeables, cette tension serait maintenue à sa valeur d'origine indéfiniment en l'absence de rayonnement ionisant. Si un rayonnement est présent, les ions agiront pour décharger partiellement la capacité et réduire la tension par rapport à sa valeur d'origine. Si une charge ΔQ (définie comme la charge positive des ions positifs ou la charge négative des électrons) est créée par le rayonnement, alors la charge totale stockée sur la capacité sera réduite de ΔQ . La tension chutera donc de sa valeur d'origine de V_0 d'une quantité ΔV donnée par :

$$\Delta V = \frac{\Delta Q}{C} \quad (10)$$

Une mesure de ΔV donne ainsi la charge d'ionisation totale ou le courant d'ionisation intégré sur la période de la mesure.

2.6 Fonctionnement des chambres d'ionisation en mode impulsionnel

2.6.1 Considérations générales

La plupart des applications des chambres d'ionisation impliquent leur utilisation en mode courant dans lequel le taux moyen de formation d'ions dans la chambre est mesuré. Comme de nombreux autres détecteurs de rayonnement, les chambres d'ionisation peuvent également fonctionner en mode impulsionnel dans lequel chaque quantum de rayonnement distinct donne lieu à une impulsion de signal distincte. Le fonctionnement en mode impulsionnel peut offrir des avantages significatifs en termes de sensibilité ou de capacité à mesurer l'énergie du rayonnement incident. Les chambres d'ionisation en mode pulsé sont utilisées dans une certaine mesure en spectroscopie de rayonnement, bien qu'elles aient été largement remplacées par des détecteurs à diodes semi-conductrices pour de nombreuses applications de ce type. Cependant, les chambres d'ionisation fonctionnant en mode impulsionnel restent des instruments importants dans certaines applications spécialisées telles que les spectromètres alpha à grande surface ou dans les détecteurs de neutrons. Le circuit équivalent d'une chambre d'ionisation fonctionnant en mode impulsionnel est représenté sur la figure (09). La tension aux bornes de la résistance de charge R est le signal électrique de base. En l'absence de toute charge d'ionisation dans la chambre d'ionisation, cette tension est nulle, et toute la tension appliquée V_0 apparaît à travers la chambre d'ionisation elle-même. Lorsqu'une particule ionisante traverse la chambre, les paires d'ions créées commencent à dériver sous l'influence du champ électrique. Comme le montrera l'analyse ci-dessous, ces charges dérivantes donnent lieu à des charges induites sur les électrodes de la chambre d'ions qui abaissent la tension de la chambre d'ions à partir de sa valeur d'équilibre V_0 .

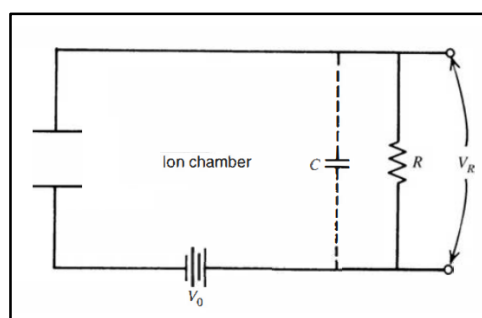


Fig. 09 Circuit équivalent d'une chambre d'ionisation fonctionnant en mode impulsionnel (C représente la capacité de la chambre plus toute la capacité parallèle V_R est la forme d'onde de l'impulsion de sortie)

Une tension apparaît alors aux bornes de la résistance de charge, qui est égale à la quantité dont la tension de la chambre a chuté. Lorsque toutes les charges à l'intérieur de la chambre ont été collectées aux électrodes opposées, cette tension atteint sa valeur maximale. Il s'ensuit alors un lent retour aux conditions d'équilibre sur une échelle de temps déterminée par la constante de temps RC

du circuit extérieur. Pendant cette période, la tension aux bornes de la résistance de charge tombe progressivement à zéro et la tension de la chambre revient à sa valeur d'origine V_0 . Si la constante de temps du circuit externe est large par rapport au temps nécessaire pour collecter les charges à l'intérieur de la chambre, une impulsion est produite dont l'amplitude reflète la grandeur de la charge totale d'origine générée dans la chambre d'ionisation.

Les temps typiques nécessaires pour collecter des électrons libres sur plusieurs centimètres sont de quelques microsecondes à cause des mobilités des électrons et des ions dans les gaz. D'autre part, les ions (positifs ou négatifs) dérivent beaucoup plus lentement et nécessitent généralement des temps de collecte de l'ordre de milliseconde. Par conséquent, si un signal qui reflète avec précision la contribution totale des ions et des électrons doit être généré, la constante de temps du circuit de collecte et les constantes de temps de formation d'impulsions ultérieures doivent être longues par rapport à une milliseconde. Dans ces conditions, la chambre d'ionisation doit être limitée à des fréquences d'impulsions très faibles pour éviter une accumulation excessive d'impulsions. De plus, la sensibilité de la sortie des circuits de mise en forme aux basses fréquences rend le système sensible aux interférences des signaux microphoniques générés par les vibrations mécaniques à l'intérieur de la chambre d'ionisation.

Pour ces raisons, les chambres d'ionisation à impulsions fonctionnent plus souvent en mode sensible aux électrons. On choisit ici une constante de temps intermédiaire entre le temps de collecte des électrons et le temps de collecte des ions. L'amplitude de l'impulsion produite ne reflète alors que la dérive des électrons et aura des temps de montée et de descente beaucoup plus rapides. Des constantes de temps de mise en forme plus courtes et des cadences beaucoup plus élevées peuvent donc être tolérées. Un sacrifice important a été fait, cependant, est que l'amplitude de l'impulsion de sortie devient maintenant sensible à la position de l'interaction de rayonnement d'origine dans la chambre et ne reflète plus uniquement le nombre total de paires d'ions formées. L'utilisation des chambres maillées plus complexes, peut surmonter cet inconvénient dans une large mesure. Le gaz de remplissage de toute chambre ionique sensible aux électrons doit bien entendu être choisi parmi les gaz dans lesquels les électrons restent sous forme d'électrons libres et ne forment pas d'ions négatifs.

2.6.2 Dérivation de la forme d'impulsion

La forme de l'impulsion produite à partir du flux de charges dans les détecteurs de géométrie arbitraire est mieux traitée en utilisant le théorème de Shockley - Ramo. Cependant, dans le cas d'une chambre d'ionisation avec des électrodes à plaques parallèles séparées par un petit espacement par rapport aux dimensions des plaques, les effets de bord peuvent être négligés et l'analyse peut être supposée ne fait intervenir qu'une seule variable : la position des charges dans la direction perpendiculaire aux plaques. Il est instructif de suivre ensuite une dérivation simplifiée de la forme

d'impulsion qui invoque des arguments basés uniquement sur la conservation de l'énergie. Cette dérivation donne des résultats corrects et évite la complexité mathématique qui pourrait obscurcir une partie du comportement physique qui est important pour la détermination des caractéristiques de l'impulsion de sortie attendue dans ces conditions.

Dans la dérivation qui suit, nous supposons qu'un champ électrique suffisant est appliqué pour que la recombinaison électron-ion soit insignifiante et aussi que les charges négatives restent sous forme d'électrons libres. Nous dérivons d'abord une expression de la forme de l'impulsion pour le cas où la constante de temps du circuit collecteur est beaucoup plus grande que le temps de collecte des ions et des électrons.

La forme de l'impulsion dépend de la configuration du champ électrique et de la position à laquelle les paires d'ions sont formées par rapport aux surfaces équipotentielles qui caractérisent la géométrie du champ. Pour simplifier l'analyse suivante, on suppose que les électrodes de la chambre sont des plaques parallèles, pour lesquelles les surfaces équipotentielles sont des plans régulièrement espacés parallèles aux surfaces des électrodes, et l'intensité du champ électrique constant est donnée par :

$$\mathcal{E} = \frac{V}{d} \quad (11)$$

V est la tension aux bornes des électrodes de la chambre et d est leur espacement. Comme autre simplification, nous supposons que toutes les paires d'ions sont formées à une distance égale x de l'électrode positive où le potentiel électrique est égal à $\mathcal{E}x$. Cette situation est donnée dans la figure (10).

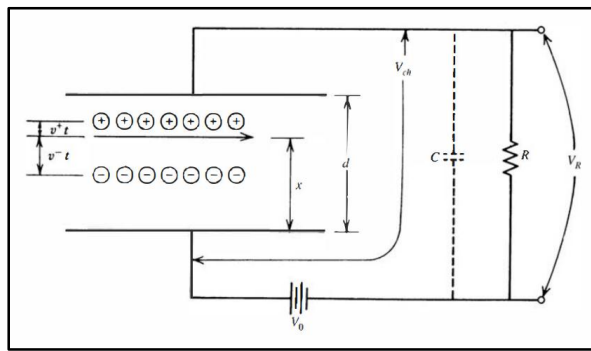


Fig. 10 Schéma de dérivation de la forme d'impulsion $V_R(t)$ pour le signal d'une chambre d'ionisation

La forme d'impulsion est plus facilement dérivée sur la base d'arguments impliquant la conservation de l'énergie. Etant donné que la constante de temps du circuit externe est supposée être grande, aucun courant appréciable ne peut circuler pendant le temps relativement court nécessaire pour collecter les

charges à l'intérieur de la chambre d'ionisation. La tension du signal est mesurée aux bornes de R sur la figure (10) et est notée V_R .

$$V_R = \frac{n_0 e}{C} \quad (12)$$

La forme de l'impulsion de signal est illustrée à la figure (11). Lorsque la constante de temps du circuit de collecte est très grande, l'amplitude maximale de l'impulsion du signal est donnée par :

$$V_{\max} = \frac{n_0 e}{C} \quad (13)$$

En fonctionnement sensible aux électrons, la partie de l'impulsion dérivée de la figure (11) qui correspond à la dérive des ions est presque entièrement perdue en choisissant une constante de temps de collecte beaucoup plus courte que le temps de collecte des ions. L'impulsion qui subsiste alors ne reflète que la dérive des électrons.

$$V_{\text{elec}} = \frac{n_0 e}{C} \cdot \frac{x}{d} \quad (14)$$

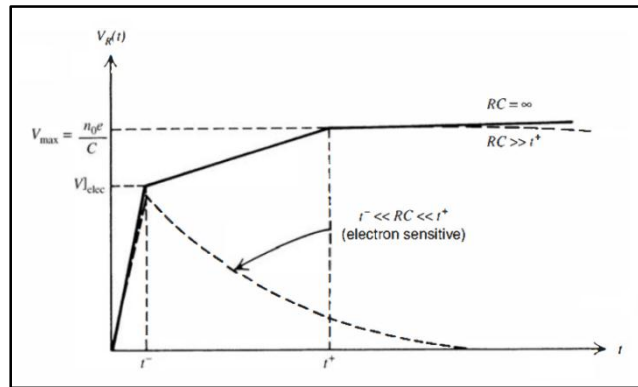


Fig. 11 Forme d'impulsion de sortie $V_R(t)$ pour différentes constantes de temps RC du schéma de principe de la figure (14)

La forme de cette impulsion est également esquissée à la figure (11). Seul le front montant rapide de l'impulsion est conservé, et l'amplitude dépend maintenant de la position x d'origine à laquelle les électrons se sont formés dans la chambre.

Dans toute situation réelle, le rayonnement incident crée des paires d'ions tout le long de son parcours à l'intérieur de la chambre. Les discontinuités nettes illustrées à la figure (11) sont alors quelque peu "délaquées" dans la forme d'impulsion résultante. Un fonctionnement sensible aux électrons conduira également à une situation dans laquelle une gamme d'amplitudes d'impulsions sera produite pour un rayonnement incident mono-énergétique est souvent un inconvénient décisif.

2.6.3 Chambre d'ionisation à grille de Frisch

La dépendance de l'amplitude de l'impulsion de la position de l'interaction dans les chambres à ions sensibles aux électrons peut être supprimée grâce à l'utilisation d'un arrangement illustré à la figure (12). Le volume de la chambre d'ionisation est divisé en deux parties par une grille de Frisch, du nom de l'auteur de la conception. Grâce à l'utilisation de la collimation externe ou de l'emplacement préférentiel de la source de rayonnement, toutes les interactions de rayonnement sont confinées dans le volume entre la grille et la cathode de la chambre. Les ions positifs dérivent simplement de ce volume vers la cathode. La grille est maintenue à un potentiel intermédiaire entre les deux électrodes et est rendue la plus transparente possible aux électrons. Les électrons sont donc attirés initialement du volume d'interaction vers la grille. En raison de l'emplacement de la résistance de charge dans le circuit, ni la dérive descendante des ions ni la dérive ascendante des électrons jusqu'à la grille ne produit de tension à mesurée. Cependant, une fois que les électrons traversent la grille pour se rendre à l'anode, la tension grille-anode commence à chuter et une tension commence à se développer à travers la résistance. La tension maximale est donc :

$$V_{\max} = \frac{n_0 e}{C} \quad (15)$$

Cette équation est identique à l'équation (13). Cependant, la tension est résultante uniquement de la dérive des électrons plutôt que du mouvement des électrons et des ions positifs. La montée lente correspondant à la dérive des ions est supprimée, et la constante de temps du circuit peut être donc fixée à une valeur beaucoup plus faible du mode de fonctionnement sensible aux électrons décrit dans la section précédente. Étant donné que chaque électron passe par la même différence de potentiel et contribue de manière égale à l'impulsion du signal de tension, l'amplitude de l'impulsion est désormais indépendante de la position de formation d'origine des paires d'ions et est simplement proportionnelle au nombre total de paires d'ions formées le long de la trajectoire de la particule incidente.

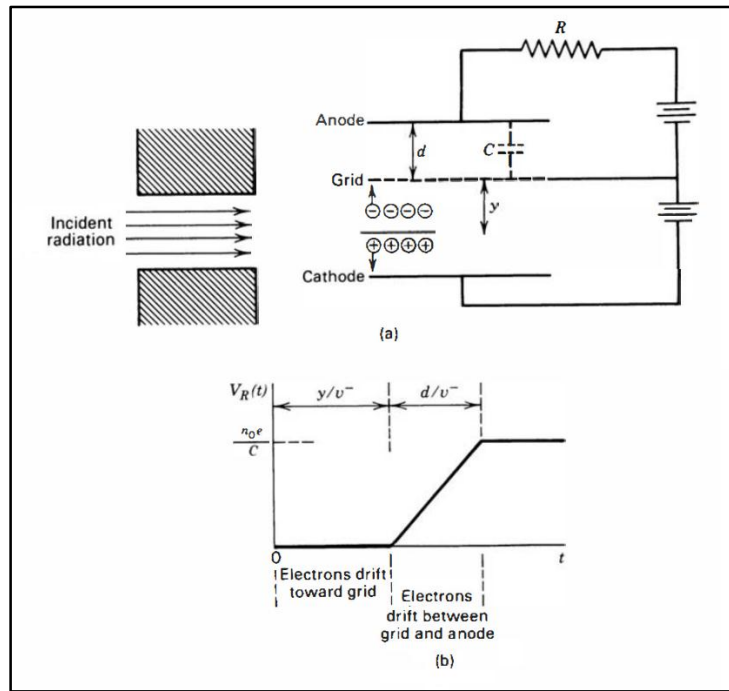


Fig. 12 (a) Principe de fonctionnement de la chambre d'ionisation à grille de Frisch
 (b) Forme d'impulsion qui résulte de la formation de n_0 paires d'ions à une distance y de la grille, où d est l'espace grille-anode

Toutes les paires d'ions doivent être formées dans la région inférieure de la chambre entre la cathode et la grille. La montée de l'impulsion résulte de la dérive des électrons à travers la région grille-anode. L'impulsion chutera à zéro avec une constante de temps donnée par RC .

Annexe I

Détecteurs à gaz (Compteur proportionnel)

1. Introduction

Le compteur proportionnel est un type de détecteur rempli de gaz qui a été introduit à la fin des années quarante. Comme les compteurs Geiger-Mueller qui seront décrits à la dernière partie de cette annexe, les compteurs proportionnels fonctionnent presque toujours en mode impulsif et s'appuient sur le phénomène de « gaz multiplication » pour amplifier la charge représentée par les paires d'ions d'origine créées dans le gaz. Les impulsions sont donc considérablement plus importantes que celles des chambres ioniques utilisées dans les mêmes conditions, et des compteurs proportionnels peuvent être appliqués à des situations dans lesquelles le nombre de paires d'ions générées par le rayonnement est trop petit pour permettre un fonctionnement satisfaisant dans des chambres d'ionisation de type impulsif. Une application importante des compteurs proportionnels a été dans la détection et la spectroscopie des rayonnements X de basse énergie (où leurs principaux concurrents sont les détecteurs à base de silicium). Les compteurs proportionnels sont également largement utilisés dans la détection des neutrons.

2. Multiplication du gaz

2.1 Formation d'avalanches

La multiplication des gaz est une conséquence de l'augmentation du champ électrique dans le gaz à une valeur suffisamment élevée. Aux faibles valeurs du champ, les électrons et les ions créés par le rayonnement incident dérivent simplement vers leurs électrodes collectrices respectives. Lors de la migration de ces charges, de nombreuses collisions se produisent normalement avec des molécules du gaz neutre. En raison de leur faible mobilité, les ions positifs ou négatifs atteignent très peu d'énergie moyenne entre les collisions.

Les électrons libres, en revanche, sont facilement accélérés par le champ appliqué et peuvent avoir une énergie cinétique importante lorsqu'ils subissent une telle collision. Si cette énergie est supérieure à l'énergie d'ionisation de la molécule du gaz neutre, il est possible qu'une paire d'ions supplémentaire soit créée lors de la collision. Étant donné que l'énergie moyenne de l'électron entre les collisions augmente avec l'augmentation du champ électrique, il existe une valeur seuil du champ au-dessus de laquelle cette ionisation secondaire se produira. Dans les gaz typiques, à pression atmosphérique, le champ de seuil est de l'ordre de 10^6 V/m.

L'électron libéré par ce processus d'ionisation secondaire sera également accéléré par le champ électrique. Lors de sa dérive ultérieure, il subit des collisions avec d'autres molécules du gaz neutre et peut ainsi créer une ionisation supplémentaire. Le processus de multiplication des gaz prend donc

la forme d'une cascade, appelée avalanche de Townsend, dans laquelle chaque électron libre créé lors d'une telle collision peut potentiellement créer plus d'électrons libres par le même processus. L'augmentation fractionnaire du nombre d'électrons par unité de longueur de la trajectoire est régie par l'équation de Townsend :

$$\boxed{\frac{dn}{n} = \alpha dx} \quad (01)$$

α est appelé le premier coefficient de Townsend du gaz. Sa valeur est nulle pour les intensités du champ électrique inférieures au seuil et augmente généralement avec l'augmentation de l'intensité du champ au-dessus de ce minimum. Pour un champ spatialement constant (comme dans la géométrie des plaques parallèles), α est une constante dans l'équation de Townsend. Sa solution prédit alors que la densité d'électrons croît de manière exponentielle avec la distance au fur et à mesure que l'avalanche progresse :

$$\boxed{n(x) = n(0)e^{\alpha x}} \quad (02)$$

Pour la géométrie cylindrique utilisée dans la plupart des compteurs proportionnels, le champ électrique augmente dans la direction dans laquelle l'avalanche progresse, et la croissance avec la distance est encore plus raide. Dans le compteur proportionnel, l'avalanche se termine lorsque tous les électrons libres ont été collectés à l'anode. Dans des conditions appropriées, le nombre d'événements d'ionisation secondaire peut être maintenu proportionnel au nombre de paires d'ions primaires formées, mais le nombre total d'ions peut être multiplié par un facteur de plusieurs milliers. Cette amplification de charge dans le détecteur lui-même réduit la nécessité d'amplificateur externe et peut entraîner des caractéristiques « signal sur bruit » considérablement améliorées par rapport aux chambres d'ionisation de type impulsif. La formation d'une avalanche implique de nombreuses collisions électron-atome énergétiques dans lesquelles une variété d'états atomiques ou moléculaires excités peut être formée. Les performances des compteurs proportionnels sont donc beaucoup plus sensibles à la composition des traces d'impuretés dans le gaz de remplissage qui n'est pas le cas pour les chambres d'ionisation.

2.2 Régions de fonctionnement du détecteur

Les différences entre les types de détecteurs à gaz fonctionnant en mode impulsif sont illustrées à la figure (01). L'amplitude de l'impulsion observée provenant du détecteur est tracée en fonction de la tension ou du champ électrique appliqué dans le détecteur. Aux très faibles valeurs de

la tension, le champ est insuffisant pour empêcher la recombinaison des paires d'ions d'origine, et la charge collectée est inférieure à celle représentée par les paires d'ions d'origine. Au fur et à mesure que la tension augmente, la recombinaison est éliminée et la région de saturation en ions est atteinte. C'est le mode de fonctionnement normal des chambres d'ionisation. Au fur et à mesure que la tension augmente encore, le champ seuil auquel commence la multiplication du gaz est atteint. La charge collectée commence alors à se multiplier et l'amplitude d'impulsion observée va augmenter. Dans une certaine plage du champ électrique, la multiplication du gaz sera linéaire et la charge collectée sera proportionnelle au nombre de paires d'ions d'origine créées par le rayonnement incident. C'est la région de la proportionnalité et représente le mode de fonctionnement des compteurs proportionnels classiques. Dans des conditions de fonctionnement constantes, l'amplitude d'impulsion observée indique toujours le nombre de paires d'ions créées dans le détecteur, bien que leur charge ait été fortement amplifiée.

L'augmentation supplémentaire de la tension ou du champ électrique appliqué peut introduire des effets non linéaires. Le plus important d'entre eux est lié aux ions positifs, qui sont également créés dans chaque processus d'ionisation secondaire. Bien que les électrons libres soient rapidement collectés, les ions positifs se déplacent beaucoup plus lentement et, pendant le temps qu'il faut pour collecter les électrons, ils se déplacent à peine. Par conséquent, chaque impulsion dans le compteur crée un nuage d'ions positifs, qui est lent à se disperser lorsqu'il dérive vers la cathode. Si la concentration de ces ions est suffisamment élevée, ils représentent une charge d'espace qui peut considérablement modifier la forme du champ électrique à l'intérieur du détecteur. Étant donné que la multiplication supplémentaire du gaz dépend de l'amplitude du champ électrique, certaines non-linéarités commenceront à être observées. Ces effets marquent le début de la région de proportionnalité limitée dans laquelle l'amplitude de l'impulsion augmente encore avec l'augmentation du nombre de paires d'ions initiales, mais pas de manière linéaire.

Si la tension appliquée est rendue suffisamment élevée, la charge d'espace créée par les ions positifs peut devenir complètement dominante dans la détermination de l'état ultérieur de l'impulsion. Dans ces conditions, l'avalanche se poursuit jusqu'à ce qu'un nombre suffisant d'ions positifs ait été créé pour réduire le champ électrique en dessous du point où une multiplication supplémentaire de gaz peut avoir lieu. Le processus est alors autolimitant et se terminera lorsque le même nombre total d'ions positifs aura été formé, quel que soit le nombre de paires d'ions initiales créées par le rayonnement incident. Alors chaque impulsion de sortie du détecteur est de même amplitude et ne reflète plus aucune propriété du rayonnement incident. Il s'agit de la région de fonctionnement Geiger-Mueller.

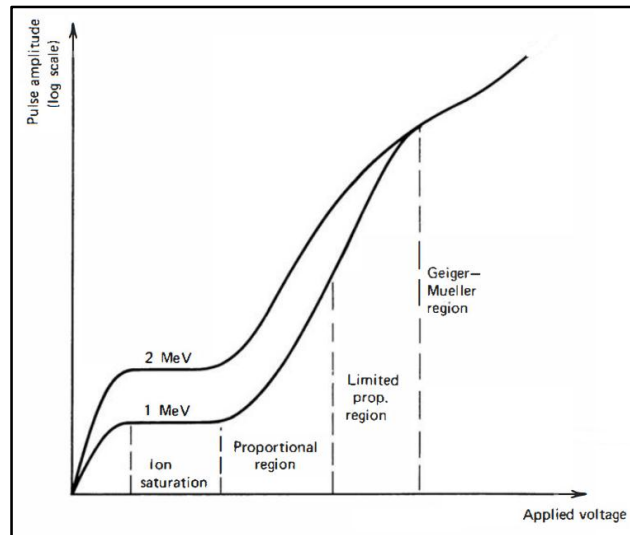


Fig. 01 Différentes régions de fonctionnement des détecteurs remplis de gaz
(L'amplitude d'impulsion observée est tracée pour des événements déposant deux quantités différentes d'énergie dans le gaz)

2.3 Choix de la géométrie

Les compteurs proportionnels typiques sont construits avec la géométrie cylindrique illustrée à la figure (02). L'anode est constituée d'un fil fin qui est positionné le long de l'axe d'un grand tube creux qui sert de cathode. La polarité de la tension appliquée dans cette configuration est importante, car les électrons doivent être attirés vers le fil axial central. Cette polarité est nécessaire de deux points de vue :

1. La multiplication des gaz nécessite de grandes intensités du champ électrique. En géométrie cylindrique, le champ électrique à un rayon r est donné par :

$$\mathcal{E}(r) = \frac{V}{r \ln(b/a)} \quad (03)$$

Où :

V = tension appliquée entre l'anode et la cathode ;

a = rayon du fil de l'anode ;

b = rayon intérieur de la cathode.

De grandes valeurs du champ électrique apparaissent au voisinage du fil de l'anode où r est petit. Comme les électrons sont attirés par l'anode, ils seront donc attirés vers la région de champ élevé.

2. Si une multiplication uniforme doit être obtenue pour toutes les paires d'ions formées par l'interaction de rayonnement d'origine, la région de multiplication du gaz doit être confinée dans un

très petit volume par rapport au volume total du gaz. Dans ces conditions, presque toutes les paires d'ions primaires sont formées en dehors de la région de multiplication, et l'électron primaire dérive simplement vers cette région avant que la multiplication n'ait lieu. Par conséquent, chaque électron subit le même processus de multiplication quelle que soit sa position d'origine de formation, et le facteur de multiplication sera le même pour toutes les paires d'ions d'origine.

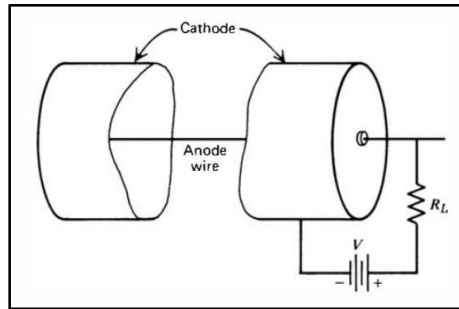


Fig. 02 Éléments de base d'un compteur proportionnel

La cathode externe doit également fournir une enceinte étanche au vide pour le gaz de remplissage et l'impulsion de sortie est développée à travers la résistance de charge R_L .

3. Caractéristiques de conception des compteurs proportionnels

3.1 Tubes scellés

Le schéma d'un compteur proportionnel incorporant de nombreuses caractéristiques de conception communes est illustré à la figure (03). L'anode à fil axial mince est supportée à chacune de ses extrémités par des isolateurs qui fournissent une liaison électrique étanche au vide pour la connexion à la haute tension. La cathode externe est classiquement mise à la terre, de sorte qu'une haute tension positive doit être appliquée pour garantir que les électrons sont attirés vers la région du champ élevé au voisinage du fil de l'anode. Pour les applications impliquant des neutrons ou des rayons gamma de haute énergie, la paroi de la cathode peut avoir plusieurs millimètres d'épaisseur pour fournir une rigidité structurelle adéquate. Pour les rayons gamma à faible énergie, les rayons X ou le rayonnement particulaire, une "fenêtre" d'entrée mince peut être prévue soit à l'une des extrémités du tube, soit à un certain point le long de la paroi de la cathode.

Une bonne résolution énergétique dans un compteur proportionnel dépend essentiellement de la garantie que chaque électron formé lors d'un événement d'ionisation originale est multiplié par le même facteur dans le processus de multiplication des gaz. L'effet mécanique le plus important qui peut bouleverser cette proportionnalité est une distorsion du champ électrique axialement uniforme prédit par l'équation (03).

D'autres causes courantes de la distorsion du champ électrique sont les effets qui se produisent près des extrémités du tube. La solution la plus courante au problème de l'effet des extrémités est de concevoir le compteur proportionnel de manière à ce que les événements se produisant près des extrémités ne subissent aucune multiplication du gaz de sorte qu'une transition abrupte se produise entre ces régions "mortes" et le reste du volume actif du tube proportionnel. D'autres solutions au problème de l'effet des extrémités impliquent la correction du champ par l'utilisation d'une plaque d'extrémité semi-conductrice ou par des anneaux conducteurs sur une plaque d'extrémité isolante qui sont maintenus au potentiel approprié pour éviter la distorsion du champ près de l'extrémité du tube.

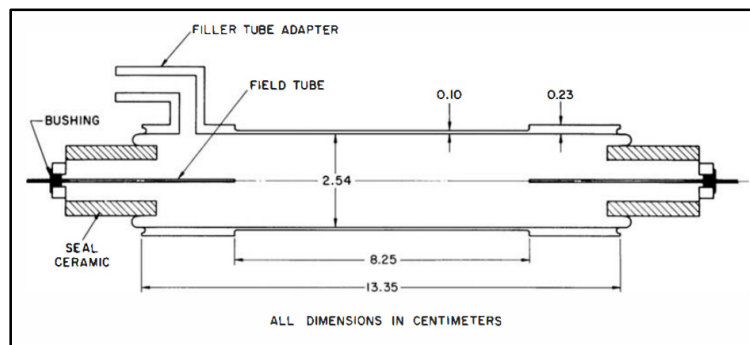


Fig. 03 Coupe transversale d'un modèle spécifique de tube proportionnel utilisé dans la détection de neutrons rapides. L'anode est un fil d'acier inoxydable de 0,025 mm de diamètre. Les tubes de champ sont constitués d'aiguilles hypodermiques de 0,25 mm de diamètre fixées autour de l'anode à chaque extrémité du tube.

Un type unique de compteur proportionnel connu sous le nom de tube de paille a été largement exploité pour la poursuite des particules en physique des hautes énergies. Les cathodes externes des tubes de petit diamètre jusqu'à plusieurs mètres de longueur sont fabriquées à partir d'une feuille métallique en utilisant un processus d'enroulement pour produire un tube autoportant similaire à une paille de soude comme illustré à la figure (04). Les fils d'anode conventionnels sont alors situés au centre du tube, et les tubes sont remplis d'un gaz proportionnel approprié. La paroi mince de la cathode est un avantage dans de nombreuses applications de poursuite pour minimiser les pertes d'énergie des particules dans les parois. Étant donné qu'un grand nombre de ces tubes sont généralement déployés pour constituer des réseaux pour des applications de poursuite, le coût relativement faible du processus de fabrication est également un avantage.

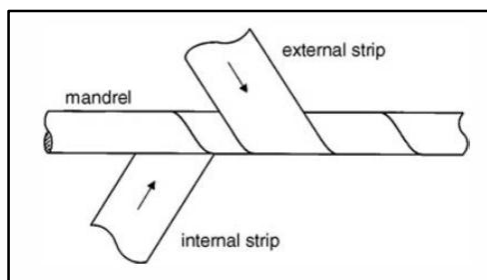


Fig. 04 Fabrication d'un tube de paille à partir de deux bandes d'aluminium
(L'adhésif est utilisé pour lier les bandes ensemble dans un tube rigide)

3.2 Compteurs de débit sans fenêtre

Une autre configuration courante du compteur proportionnel est illustrée à la figure (05). On suppose que la source de rayonnement est un petit échantillon d'un radio-isotope, qui peut ensuite être introduit directement dans le volume de comptage hémisphérique du détecteur. Le grand avantage du comptage interne de la source réside dans le fait qu'aucune fenêtre d'entrée n'est nécessaire entre la source de rayonnement et le volume actif du compteur. Étant donné que les matériaux de fenêtre peuvent sérieusement atténuer les rayonnements mous tels que les rayons X ou les particules alpha, la configuration de la source interne est la plus largement utilisée pour ces applications.

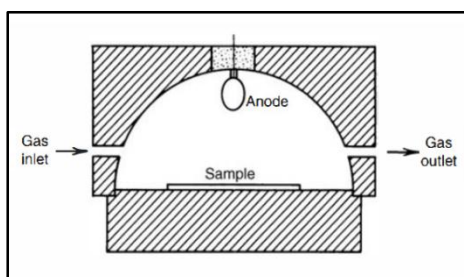


Fig. 05 Schéma d'un compteur proportionnel à débits de gaz
avec un fil d'anode en boucle et un volume hémisphérique
(L'échantillon peut souvent être inséré dans la chambre en faisant
glisser un plateau pour minimiser la quantité d'air introduit)

La géométrie de comptage peut également être rendue très favorable pour ce type de détecteur. Dans le système représenté sur la figure (05), pratiquement tout quantum émergeant de la surface de la source se retrouve dans le volume actif du compteur et peut générer une impulsion de sortie.

Une autre géométrie utile pour les compteurs proportionnels est le détecteur "pancake", illustré à la figure (06). Souvent utilisé dans les systèmes destinés à compter l'activité alpha et bêta, l'échantillon est placé près de l'une des surfaces planes du détecteur pour fournir une géométrie de comptage proche de 2π . Cette conception présente certains avantages par rapport au type hémisphérique, elle a généralement un volume plus petit et donc un fond plus faible, et limite

également la longueur des traces de particules « bêta » dans le gaz pour minimiser leur amplitude d'impulsion et ainsi maximiser la séparation en amplitude des particules alpha.

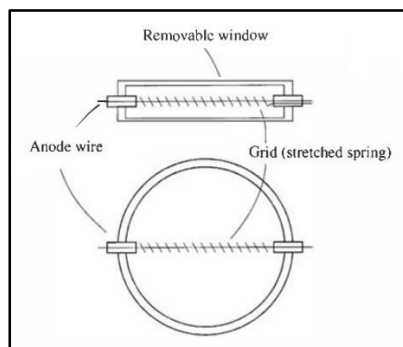


Fig. 06 Forme du tube proportionnel "pancake"

3.3 Gaz de remplissage

Étant donné que la multiplication des gaz dépend essentiellement de la migration des électrons libres plutôt que des ions négatifs beaucoup plus lents, le gaz de remplissage dans les compteurs proportionnels doit être choisi parmi ceux qui ne présentent pas de coefficient d'attachement électronique important. Comme l'air n'en fait pas partie, les compteurs proportionnels doivent être conçus de manière à maintenir la pureté du gaz. Le gaz peut être scellé de manière permanente dans le compteur ou circuler lentement à travers le volume de la chambre dans les conceptions du type à flux continu. Les compteurs étanches sont plus pratiques à utiliser, mais leur durée de vie est parfois limitée par des fuites microscopiques qui entraînent une contamination progressive du gaz de remplissage. Les compteurs à débit continu nécessitent des systèmes d'alimentation en gaz qui peuvent être encombrants mais évitent de nombreux problèmes potentiels liés à la pureté du gaz.

La multiplication des gaz dans le compteur proportionnel est basée sur l'ionisation secondaire créée lors des collisions entre les électrons et les molécules de gaz neutre. Outre l'ionisation, ces collisions peuvent également produire une simple excitation de la molécule de gaz sans création d'électron secondaire. Ces molécules excitées ne contribuent pas directement à l'avalanche mais se désintègrent à leur état fondamental par l'émission d'un photon visible ou ultraviolet. Dans les circonstances appropriées, ces photons de désexcitation pourraient créer une ionisation supplémentaire ailleurs dans le gaz de remplissage par des interactions photoélectriques avec des couches d'électrons moins étroitement liées ou pourraient produire des électrons par des interactions au niveau de la paroi du compteur. Bien que de tels événements induits par des photons soient importants dans la région de fonctionnement de Geiger-Mueller, ils sont généralement indésirables dans les compteurs proportionnels car ils peuvent entraîner une perte de proportionnalité et/ou générer des impulsions parasites.

4. Performance du compteur proportionnel

4.1 Facteur de multiplication de gaz

Une étude du processus de multiplication dans les gaz est normalement divisée en deux parties. La réponse électronique unique du compteur est définie comme la charge totale qui est développée par la multiplication gazeuse si l'avalanche est initiée par un seul électron provenant de l'extérieur de la région de multiplication gazeuse. Ce processus peut être étudié expérimentalement en créant des conditions d'irradiation dans lesquelles un seul électron est libéré par interaction (par exemple, par l'interaction photoélectrique des photons ultraviolets incidents). Des facteurs de multiplication de l'ordre de 10^5 sont suffisants pour permettre une détection directe des impulsions résultantes, et leur répartition en amplitude peut fournir des informations sur les mécanismes de multiplication des gaz au sein du compteur. Les résultats de ces expériences sont utilisés pour étudier la résolution en énergie des compteurs proportionnels.

Si la réponse à un seul électron est connue, les propriétés de l'amplitude des impulsions produites par de nombreuses paires d'ions d'origine peuvent être déduites. À condition que les effets de charge d'espace ne soient pas assez importants pour déformer le champ électrique, chaque avalanche est indépendante, et la charge totale Q générée par n_0 paires d'ions d'origine est :

$$Q = n_0 e M \quad (04)$$

Où M est le facteur moyen de multiplication du gaz qui caractérise l'opération du compteur.

Diverses expressions analytiques qui apparaissent dans la littérature pour relier le facteur de multiplication des gaz à des paramètres expérimentaux diffèrent les uns des autres par la forme de l'expression utilisée pour $\alpha(\mathcal{E})$ où le coefficient α est une fonction du type de gaz et de l'amplitude du champ électrique $\mathcal{E}(r)$.

En supposant une linéarité entre α et $\mathcal{E}$, une expression largement utilisée pour M est :

$$\ln M = \frac{V}{\ln(b/a)} \cdot \frac{\ln 2}{\Delta V} \left(\ln \frac{V}{pa \ln(b/a)} - \ln K \right) \quad (05)$$

Où :

M = facteur multiplicateur du gaz ;

V = tension appliquée ;

a = rayon de l'anode ;

b = rayon cathodique ;

p = pression du gaz.

Selon le modèle utilisé dans la dérivation, ΔV correspond à la différence de potentiel à travers laquelle un électron se déplace entre des événements ionisants successifs, et K représente la valeur minimale de $\mathcal{E}/p$ en dessous de laquelle la multiplication ne peut pas se produire. ΔV et K doivent être constantes pour tout gaz de remplissage donné.

Pour un tube proportionnel donné à pression de gaz constante, l'équation (05) montre que la multiplication du gaz augmente rapidement avec la tension appliquée V . En négligeant le terme logarithmique variant lentement, la multiplication M varie principalement comme une fonction exponentielle de V . Les compteurs proportionnels doivent donc fonctionner avec des alimentations en tension extrêmement stables pour empêcher les variations de M au cours de la mesure.

4.2 Effets des charges d'espace

Dans le processus d'avalanche dont les compteurs proportionnels dépendent, des électrons et des ions positifs sont créés. Les électrons sont collectés relativement rapidement (en quelques nanosecondes) à l'anode, laissant derrière eux les ions positifs qui se déplacent beaucoup plus lentement et diffusent progressivement en dérivant vers l'extérieur vers la paroi de la cathode du tube. La charge d'espace représentée par ces charges positives peut, dans certaines circonstances, changer sensiblement l'intensité du champ électrique par rapport à sa valeur sans charge d'espace. Étant donné que les ions se forment préférentiellement près du fil d'anode où se produit la plupart de la multiplication du gaz, l'effet de la charge d'espace sera de réduire le champ électrique à de petits rayons en dessous de sa valeur normale.

Il existe deux catégories différentes d'effets de charge d'espace. Les effets auto-induits qui surviennent lorsque le gain du gaz est suffisamment élevé pour que les ions positifs formés au cours d'une avalanche donnée puissent modifier le champ et réduire le nombre d'électrons produits dans les étapes ultérieures de la même avalanche. Cet effet dépend du degré de la multiplication du gaz et de la géométrie du tube mais ne dépend pas de la fréquence des impulsions. L'effet de charge d'espace général comprend l'effet cumulatif des ions positifs créés à partir de nombreuses avalanches différentes. Cet effet peut être important à des valeurs inférieures de la multiplication des gaz et devient plus grave à mesure que le taux d'événements à l'intérieur du tube augmente.

4.3 Résolution énergétique

4.3.1 Considérations statistiques

La charge Q , qui est développée dans une impulsion à partir d'un compteur proportionnel en l'absence d'effets non linéaires, peut être supposée être la somme des charges créées dans chaque avalanche individuelle. Il y aura n_0 de ces avalanches, chacune déclenchée par un électron distinct des n_0 paires d'ions créées par le rayonnement incident. Dans ce qui suit, nous laissons A représenter le facteur de multiplication d'électrons pour toute avalanche déclenchée par un seul électron et M représenter le facteur de multiplication moyen de toutes les avalanches qui contribuent à une impulsion donnée :

$$M = \frac{1}{n_0} \sum_{i=1}^{n_0} A_i \equiv \bar{A} \quad (06)$$

De plus,

$$M = \frac{Q}{en_0} \quad (07)$$

Ou bien

$$Q = n_0 e M \quad (08)$$

L'amplitude d'impulsion du détecteur est proportionnelle à Q . Cette amplitude est sujette à des fluctuations d'une impulsion à l'autre même dans le cas d'un dépôt d'énergie égale par le rayonnement incident car à la fois n_0 et M dans l'équation (08) montreront une certaine variation inhérente. Comme ces facteurs sont supposés indépendants, nous pouvons utiliser la formule de propagation d'erreur pour prédire la variance relative attendue de Q :

$$\left(\frac{\sigma_Q}{Q}\right)^2 = \left(\frac{\sigma_{n_0}}{n_0}\right)^2 + \frac{1}{n_0} \left(\frac{\sigma_A}{\bar{A}}\right)^2 \quad (09)$$

Cette expression permet une analyse du taux de la variance prédit de l'amplitude de l'impulsion $(\sigma_Q/Q)^2$ en termes de contributions séparées des fluctuations de paires d'ions $(\sigma_{n_0}/n_0)^2$ et des variations de multiplication d'un seul électron $(\sigma_A/\bar{A})^2$.

4.3.1.1 Variations du nombre de paires d'ions :

Le premier terme de l'équation (09) représente la variance relative du nombre original n_0 de paires d'ions. Ces fluctuations peuvent être exprimées en termes de facteur de Fano F :

$$\left(\frac{\sigma_{n_0}}{n_0} \right)^2 = \frac{F}{n_0} \quad (10)$$

Les estimations de la grandeur du facteur de Fano pour les gaz proportionnels vont d'environ 0,05 à 0,20, et certaines valeurs spécifiques sont données par la suite au tableau (01). Les valeurs les plus basses de F sont généralement observées pour les mélanges de gaz binaires.

Tab. 01 : Constantes liées à la résolution des gaz proportionnels

Resolution-Related Constants for Proportional Gases						
Gas	W (eV/ion pair)	Fano Factor F		Multiplication Variance b	Energy Resolution at 5.9 keV	
		Calculated	Measured		Calculated ^a	Measured
Ne	36.2	0.17		0.45	14.5%	
Ar	26.2	0.17		0.50	12.8%	
Xe	21.5		≤ 0.17			
Ne + 0.5% Ar	25.3	0.05		0.38	10.1%	11.6%
Ar + 0.5% C ₂ H ₂	20.3	0.075	≤ 0.09	0.43	9.8%	12.2%
Ar + 0.8% CH ₄	26.0	0.17	≤ 0.19			
Ar + 10% CH ₄	26			0.50	12.8%	13.2%
^a Given by $2.35[W(F + b)/5900 \text{ eV}]^{1/2}$ [see Eq. (27)]						

4.3.1.2 Variations dans les avalanches à un seul électron

Le deuxième terme de l'équation (09) représente la contribution des fluctuations de l'amplitude de l'avalanche à électron unique.

Une prédiction théorique simple pour la distribution d'avalanche à un seul électron peut être effectuée en supposant que la probabilité d'ionisation par un électron ne dépend que de l'intensité du champ électrique et est indépendante de son historique antérieur. On peut alors montrer que la distribution attendue du nombre d'électrons produits dans une avalanche donnée doit être prédite par la distribution de Furry,

$$P(A) = \frac{(1 - 1/\bar{A})^{A-1}}{\bar{A}} \quad (11)$$

Où :

A : multiplication d'avalanche, ou nombre d'électrons dans l'avalanche ;

$\bar{A}$: valeur moyenne de : A ($= M$).

Si A est suffisamment grand, supérieur à 50 ou 100 comme c'est presque toujours le cas, alors la distribution de « Furry » se réduit à une simple forme exponentielle,

$$P(A) \cong \frac{e^{-A/\bar{A}}}{\bar{A}} \quad (12)$$

Qui prédit une variance relative :

$$\left(\frac{\sigma_A}{\bar{A}}\right)^2 = 1 \quad (13)$$

Des expériences menées avec de faibles valeurs du champ électrique tendent à confirmer la forme exponentielle prédite de la distribution d'avalanche. Cependant, à des champs plus élevés, qui sont plus typiques pour l'utilisation du compteur proportionnel, un modèle un peu plus complexe doit être utilisé pour représenter de manière adéquate les résultats expérimentaux.

4.3.1.3 Limite statistique globale

Indépendamment de la forme supposée pour la distribution d'avalanche à électron unique, la distribution en amplitude d'impulsion Q pour de grandes valeurs de n_0 (par exemple $n_0 > 20$) se rapproche d'une distribution en forme Gaussienne. Puisque cela est vrai dans la plupart des applications, un pic de forme gaussienne symétrique dans la distribution d'amplitude d'impulsion du rayonnement mono-énergétique doit être attendu à partir des fluctuations de n_0 et A.

Afin de prédire la variance relative de cette distribution, nous revenons à l'équation (09) et en évaluant $(\sigma_{n_0}/n_0)^2$ et $(\sigma_A/\bar{A})^2$ en fonction de F et b , on obtient :

$$\begin{aligned} \left(\frac{\sigma_Q}{Q}\right)^2 &= \frac{F}{n_0} + \frac{b}{n_0} \\ \left(\frac{\sigma_Q}{Q}\right)^2 &= \frac{1}{n_0}(F + b) \end{aligned} \quad (14)$$

Où F est le facteur Fano (valeur typique de 0,05 à 0,20) et b est le paramètre de la distribution « Polya » qui caractérise les statistiques d'avalanche (valeur typique de 0,4 à 0,7). D'après les amplitudes relatives de F et b, nous voyons que la variance de l'amplitude des impulsions est dominée par les fluctuations de la taille de l'avalanche et que les fluctuations du nombre initial de paires d'ions sont généralement un faible facteur contributif.

L'écart type relatif de la distribution d'amplitude d'impulsion est obtenu en prenant la racine carrée de l'équation (14) :

$$\boxed{\frac{\sigma_Q}{Q} = \sqrt{\frac{F+b}{n_0}}} \quad (15)$$

Puisque $n_0 = E/W$, où E est l'énergie déposée par le rayonnement incident et W est l'énergie nécessaire pour former une paire d'ions, l'équation (15) devient :

$$\boxed{\frac{\sigma_Q}{Q} = \sqrt{\frac{W(F+b)}{E}}} \quad (16)$$

Où le terme $W(F+b)$ est constant pour un gaz de remplissage donné. On s'attend donc à ce que la limite statistique de la résolution en énergie d'un compteur proportionnel varie en sens inverse de la racine carrée de l'énergie déposée par le rayonnement incident.

La résolution limite est proportionnelle à $W(F+b)$, ce paramètre peut donc servir de guide lors de la comparaison de la résolution potentielle pouvant être obtenue à partir de différents gaz. Les mesures effectuées à l'aide des compteurs proportionnels à gaz courants confirment que les meilleures résolutions en énergie sont obtenues en utilisant de faibles valeurs de la tension appliquée et également en choisissant le fil d'anode de plus petit diamètre compatible avec une surface de bonne qualité avec une haute uniformité.

4.3.2 Autres facteurs affectant la résolution énergétique

Dans des circonstances bien contrôlées, la résolution énergétique réellement observée pour les compteurs proportionnels est très proche de la limite statistique prédite par l'équation (16). Certaines données concernant la résolution en énergie observées d'un compteur proportionnel pour les photons de basse énergie sont présentées à la figure (09). Afin d'atteindre des résolutions énergétiques qui se rapprochent de cette limite statistique, il faut veiller à minimiser les effets potentiellement nocifs du bruit électronique, des non-uniformités géométriques dans la chambre et des variations des paramètres de fonctionnement du détecteur.

Le plus critique des facteurs géométriques est l'uniformité et la douceur du fil d'anode utilisé dans le compteur proportionnel. Des variations du diamètre du fil aussi petites que 0,5 % peuvent être significatives. D'autres facteurs incluent l'excentricité possible du fil ou la non-uniformité de la cathode cylindrique, mais ce sont des facteurs nettement moins critiques. Les extrémités de la chambre présentent aussi des problèmes particuliers, à moins que l'on prenne soin de façonner le champ électrique en utilisant des tubes de champ ou par d'autres méthodes.

Les paramètres de fonctionnement qui peuvent affecter la résolution énergétique sont la pureté du gaz, la pression du gaz et la stabilité de la haute tension appliquée à la chambre. Des traces de gaz électronégatifs peuvent réduire considérablement la multiplication des gaz et introduire des fluctuations de l'amplitude des impulsions. Des changements de pression de gaz de l'ordre de quelques dixièmes pour cent peuvent être significatifs, comme en témoigne le changement correspondant de multiplication de gaz.

Ceci est rarement un problème pour les tubes scellés (à l'exception de la flexion de la fenêtre causée par les changements de pression barométrique) mais peut être une cause importante de perte de résolution pour les compteurs proportionnels de débit de gaz dans lesquels la pression de gaz n'est pas bien stabilisée. Enfin, l'extrême dépendance de la multiplication de gaz de la tension appliquée est indiquée par l'équation (10). Des variations de la tension appliquée de 0,1 ou 0,2 % peuvent modifier considérablement le facteur de multiplication du gaz dans des conditions typiques, et donc les alimentations en tension doivent être bien stabilisées et exemptes de dérives à long terme si l'ultime résolution énergétique d'un compteur proportionnel doit être préservée.

La dépendance du gain de ces paramètres peut être prédite en utilisant l'un des modèles physiques pour le processus de multiplication de gaz, tel que celui conduisant à l'équation (05). Quelques résultats typiques sont donnés dans le tableau (02).

Tab. 02 : Variations de gain (hauteur d'impulsion) dans un compteur proportionnel

Gain (Pulse Height) Variations in a Proportional Counter	
Predictions of the Diethorn model for a P-10 filled proportional tube, $a = 0.03$ cm, $b = 1$ cm, $V = 1793$ V, $M = 1000$:	
1 % increase in	Change in M:
Applied voltage V	+17.4%
Gas pressure p	-8.6%
Anode radius a	-6.1%
Cathode radius b	-2.7%

Il est aussi couramment observé que les changements du taux de comptage des compteurs proportionnels peuvent influencer le facteur de multiplication de gaz. Si le taux d'irradiation varie au cours d'une mesure, le changement de gain associé peut affecter négativement la résolution en énergie. Quelle que soit la ou les causes spécifiques, l'effet du taux de comptage s'est être avéré une source importante de perte de résolution dans de nombreuses applications de compteurs proportionnels.

4.4 Déficit balistique

Le temps de montée de l'impulsion du préamplificateur correspond normalement au temps de collecte de charge dans le détecteur lui-même. Si l'amplitude complète de l'impulsion du préamplificateur doit être préservée pendant le processus de mise en forme, les constantes de temps de mise en forme doivent être importantes par rapport au temps de montée de l'impulsion du préamplificateur. Etant donné que les constantes de temps de mise en forme ne peuvent pas être toujours choisies arbitrairement grandes, l'amplitude de l'impulsion mise en forme peut être parfois légèrement inférieure à celle pouvant être atteinte avec des constantes de temps très longues. Le degré auquel l'amplitude de la constante de temps infinie a été diminuée par le processus de mise en forme est appelé déficit balistique.

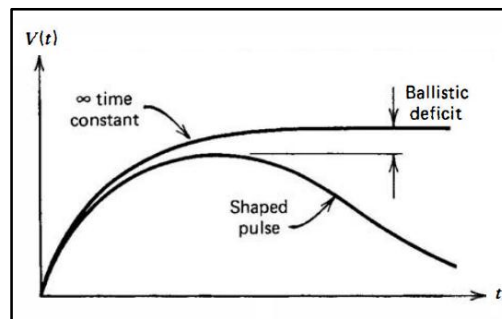


Fig. 07 Définition du déficit balistique.

Si les temps de mise en forme sont fixes, le rapport du déficit sur l'amplitude de l'impulsion mise en forme sera constant pour des impulsions de même forme de front montant mais variera si cette forme change. Dans les détecteurs à temps de collecte de charge constant, des déficits balistiques relativement importants peuvent être souvent tolérés car un taux constant de l'amplitude de chaque impulsion est perdu. Cependant, si le temps de collecte de charge varie, une quantité variable de chaque impulsion sera perdue, ce qui peut entraîner une dégradation de la résolution. Dans ces cas, il faut alors choisir des constantes de temps plus longues que celles qui pourraient être optimales sur la base des considérations du rapport signal/bruit ou d'empilement.

4.5 Caractéristiques temporelles de l'impulsion du signal

Le même type d'analyse effectuée sur les chambres ioniques à impulsions peut être appliqué pour dériver la forme de l'impulsion de sortie à partir de compteurs proportionnels. Cependant, plusieurs différences majeures résultant des considérations physiques suivantes modifient le caractère de l'impulsion de sortie.

1. Pratiquement toute la charge générée dans le tube proportionnel provient de la région d'avalanche, quel que soit l'endroit où les paires d'ions d'origine sont formées. L'historique temporelle de

l'impulsion de sortie est donc commodément divisée en deux étapes : le temps de dérive nécessaire pour que les électrons libres créés par le rayonnement se déplacent de leur position d'origine vers la région proche du fil d'anode où la multiplication peut avoir lieu, et le temps de multiplication nécessaire depuis le début de l'avalanche jusqu'à son achèvement.

Pour toute paire d'ions d'origine donnée, la contribution à l'impulsion de sortie pendant le temps de dérive est négligeable par rapport à la contribution du nombre beaucoup plus grand de porteurs de charge formés dans l'avalanche suivante. L'effet du temps de dérive est donc d'introduire un retard entre le temps de formation de la paire d'ions et le début de l'impulsion de sortie correspondante. Le temps de dérive (peut-être quelques microsecondes ou plus) est normalement bien supérieur au temps de multiplication et varie en fonction de la position radiale de la paire d'ions d'origine dans le tube.

2. Étant donné que la plupart des ions et des électrons sont créés très près du fil d'anode, la majeure partie de l'impulsion de sortie est attribuable à la dérive des ions positifs plutôt qu'au mouvement des électrons. Au début, les ions positifs se trouvent dans une région à champ élevé et se déplacent rapidement, ce qui entraîne une composante initiale de l'impulsion à montée rapide. Finalement, cependant, les ions atteignent des régions du tube à des rayons plus grands où le champ est plus petit, et leur vitesse de dérive diminue. La dernière partie de l'impulsion monte donc très lentement et n'est pas souvent observée en pratique en raison des temps de mise en forme finis des circuits électroniques. La figure (08) montre les étapes du développement d'une impulsion déclenchée par une seule particule chargée.

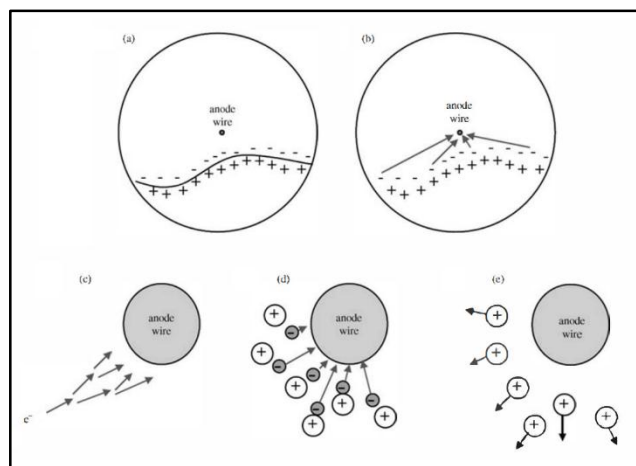


Fig. 08 Séquence d'événements dans la génération d'une impulsion dans un compteur proportionnel

- (a) Un électron énergétique ou une particule chargée produit des paires d'ions lors de son passage à travers le compteur à gaz (nanoseconde).

- (b) Les électrons d'ionisation dérivent à l'intérieur le long d'un rayon vers le fil d'anode (quelques microsecondes).
- (c) Chaque électron arrivant dans la région de multiplication déclenche sa propre avalanche (inférieure à la microseconde).
- (d) Les centaines ou milliers d'électrons formés dans chaque avalanche sont rapidement attirés vers la surface du fil (inférieure à la microseconde).
- (e) Les ions positifs beaucoup plus lents de l'avalanche se déplacent vers l'extérieur à travers la région de champ élevé près du fil, produisant la majeure partie de la sortie amplitude d'impulsion (quelques microsecondes).

5. Efficacité de détection et courbes de comptage

5.1 Sélection de la tension de fonctionnement

Pour les rayonnements chargés tels que les particules alpha ou bêta, une impulsion de signal sera produite pour chaque particule qui dépose une quantité importante d'énergie dans le gaz de remplissage. À des valeurs élevées de multiplication du gaz, une seule paire d'ions peut déclencher une avalanche avec une ionisation secondaire suffisante pour être détectable à l'aide de préamplificateurs de faible bruit. Sauf si l'application l'exige, les compteurs proportionnels fonctionnent rarement dans un mode sensible aux avalanches simples car la mesure est alors sujette à des non-linéarités pour des impulsions plus importantes en raison des effets de charge d'espace. De plus, de nombreuses impulsions de fond et parasites peuvent être également comptées, ce qui correspond souvent à très peu de paires d'ions d'origine.

Au lieu de cela, des valeurs inférieures de la multiplication du gaz sont généralement utilisées, ce qui nécessite que l'impulsion provienne d'un nombre fini de paires d'ions afin d'avoir une amplitude suffisamment grande pour dépasser le niveau de discrimination du système de comptage. Étant donné que le facteur de multiplication du gaz varie avec la tension appliquée au détecteur, une "courbe de comptage" peut être enregistrée pour sélectionner une tension de fonctionnement appropriée à l'application spécifique. Dans cette mesure, le taux de comptage est enregistré dans des conditions de source constante lorsque la tension du détecteur varie. Un point de fonctionnement est alors sélectionné qui normalement correspond à une région plate ou "plateau" sur la courbe résultante de la vitesse en fonction de la tension.

5.2 Comptage alpha

Si presque toutes les impulsions du détecteur sont de la même taille, le spectre différentiel d'hauteur d'impulsion a un seul pic isolé, comme l'indique la figure (09).

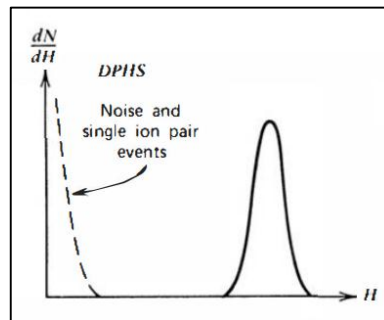


Fig. 09 Spectre différentiel d'hauteur d'impulsion d'une source alpha

La courbe de comptage correspondante représentée à la figure (10) présente alors un palier simple. Un point de fonctionnement sur le plateau de comptage assure alors que toutes les impulsions issues du détecteur sont comptées. Ce comportement simple est une caractéristique du compteur Geiger-Mueller mais n'est observé dans les compteurs proportionnels que dans des circonstances particulières. Un tel cas concerne les particules chargées mono-énergétiques dont la portée dans le compteur à gaz est inférieure aux dimensions de la chambre à gaz. Les sources de particules alpha entrent souvent dans cette catégorie, et les compteurs proportionnels peuvent donc facilement enregistrer chaque particule qui pénètre dans le volume actif avec une efficacité pratiquement de 100 %. Pour éviter les pertes d'énergie dans les fenêtres d'entrée des tubes scellés, les compteurs sans fenêtre de flux sont souvent utilisés pour le comptage alpha. Les mesures absolues de l'activité de la source alpha sont donc relativement simples et impliquent l'évaluation de l'angle solide effectif sous-tendu par le volume actif du compteur (qui peut être proche de 2π pour les compteurs de flux de source interne) et de petites corrections pour l'auto-absorption et la diffusion alpha dans la source lui-même ou dans le support source.

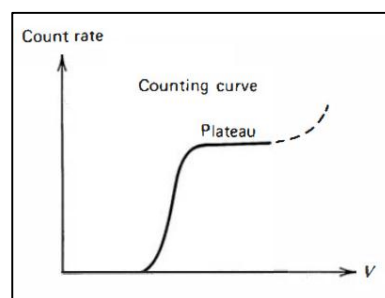


Fig. 10 Courbe de comptage d'une source alpha

5.3 Comptage beta

Pour les particules « bêta » d'énergies typiques, la portée de particules dépasse largement les dimensions de la chambre à gaz. Le nombre de paires d'ions formées dans le gaz est alors proportionnel à la seule petite fraction de l'énergie des particules dissoute dans le gaz avant d'atteindre la paroi opposée. Les impulsions des particules « bêta » sont donc normalement beaucoup plus petites que celles induites par les particules alpha d'énergie cinétique équivalente et couvriront également une plage d'amplitude plus large en raison de la propagation des énergies des particules « bêta » et les variations des trajets possibles dans le gaz. Un spectre de hauteur d'impulsion différentielle, qui pourrait représenter une source mixte typique d'activité bêta et alpha, est schématisé par la figure (11).

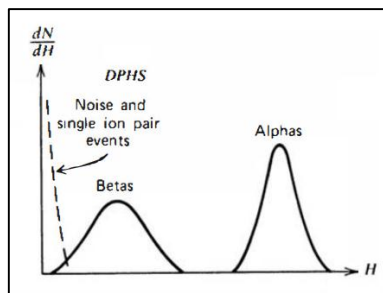


Fig. 11 Spectre différentiel d'impulsion d'une source mixte bêta et alpha

La courbe de comptage correspondante montre maintenant deux plateaux (voir figure 12) : le premier au point où seules les particules alpha sont comptées, et le second où les particules alpha et bêta sont comptées. Étant donné que la distribution de la hauteur des impulsions des particules « bêta » est plus large et moins bien séparée du bruit de faible amplitude, le "plateau bêta" est généralement plus court et présente une pente plus importante que le "plateau alpha".

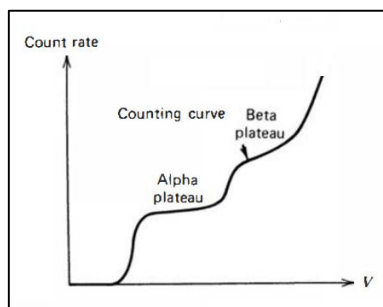


Fig. 12 Courbe de comptage d'une source mixte bêta et alpha

5.4 Sources mixtes

Dans de nombreuses applications, l'objectif sera de compter séparément les activités alpha et bêta qui peuvent être mixées dans un échantillon commun. La séparation inhérente entre les amplitudes attendues des événements induits par alpha et bêta, (figure 11), peut être exploitée pour fournir deux canaux de lecture séparés, l'un correspondant au compte alpha et l'autre au compte bêta. Cette sélection est accomplie en définissant électroniquement des fenêtres de hauteur d'impulsion dans des branches parallèles du circuit de comptage électronique, une fenêtre définie pour englober le pic alpha et l'autre pour inclure les particules bêta. Si l'uniformité du gain et la résolution en énergie de la chambre proportionnelle sont suffisamment bonnes et que sa géométrie est telle que les deux rayonnements contribuent à des impulsions d'amplitude très différente, il en résultera une séparation nette et il y aura très peu de "diaphonie" entre les deux voies de comptage. De mauvaises géométries qui permettent aux particules bêta de parcourir de longues distances dans le gaz ou qui introduisent une variabilité dans les longueurs de piste alpha réduira la séparation des événements et peut induire une diaphonie importante si les amplitudes se chevauchent. Cette diaphonie peut être réduite ou éliminée en rétrécissant les fenêtres d'acceptation de hauteur d'impulsion, mais uniquement au prix des efficacités de comptage réduites.

L'un des avantages des compteurs proportionnels est que la séparation alpha/bêta inhérente en amplitude se traduit par un taux de comptage de fond très faible dans le canal alpha. La plupart des sources de fond telles que les rayons cosmiques ou les rayons gamma ambiants, produisent une ionisation de faible densité des particules bêta et ont donc une amplitude trop faible pour tomber dans le canal alpha. Le bruit de fond résiduel est en grande partie dû à l'activité alpha naturelle du matériau de construction du compteur et peut être aussi faible que 1 ou 2 coups par heure.

5.5 Sources de rayons X et de rayons gamma

Les compteurs proportionnels peuvent être utilisés pour la détection et la spectroscopie des rayons X mous ou des rayons gamma dont l'énergie est suffisamment faible pour interagir avec une efficacité raisonnable avec le gaz du compteur. La spectroscopie des rayons X de basse énergie est l'une des applications les plus importantes des compteurs proportionnels et repose sur l'absorption totale des photoélectrons formés par interactions de photons dans le gaz. Étant donné que l'énergie des photoélectrons est directement liée à l'énergie des rayons X, les énergies des photons peuvent être identifiées à partir de la position des pics correspondants de "pleine énergie" dans le spectre d'hauteur d'impulsion mesuré.

Les rayons X de basse énergie produisent des photoélectrons avec des énergies minimales de quelques dizaines de ke V. Avec des valeurs W typiques de 25 à 35 eV/paire d'ions, seules quelques centaines de paires d'ions sont créées dans le gaz le long de la trajectoire d'un électron pour servir de base à l'impulsion du signal. La charge électrique correspondante n'est pas beaucoup plus grande que la charge de bruit équivalente à l'entrée

des préamplificateurs à faible bruit. Ainsi, des gains "internes" supplémentaires grâce au processus d'amplification de gaz inhérent aux compteurs proportionnels sont essentiels pour améliorer le rapport signal sur bruit des impulsions mesurées.

La fonction de réponse pour les rayons X de faible énergie peut être compliquée par plusieurs effets liés aux rayons X caractéristiques générés par l'interaction du rayonnement primaire dans le détecteur.

Le plus important concerne les rayons X caractéristiques K qui suivent généralement l'absorption photoélectrique du rayonnement primaire dans le gaz de remplissage. Puisque l'énergie correspondante peut être relativement grande, ce rayon X peut s'échapper du gaz sans autre interaction. Comme le montre la figure (13), un "pic d'échappement" correspondant apparaîtra alors dans la fonction de réponse qui se situe en dessous du pic de pleine énergie d'une quantité égale à l'énergie caractéristique des rayons X. Ce pic d'échappement peut être supprimé en choisissant un gaz dont l'énergie de liaison de la couche K se situe au-dessus de l'énergie des rayons X incidents. D'autres pics de la fonction de réponse peuvent également provenir de l'absorption des rayons X caractéristiques générés par l'interaction du rayonnement primaire dans la fenêtre d'entrée ou les parois du compteur. Un matériau de fenêtre à faible Z tel que le béryllium est souvent choisi pour minimiser cette contribution.

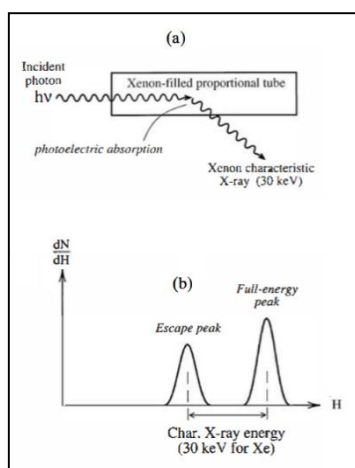


Fig. 13 La partie (a) montre le processus qui donne lieu au pic d'échappement des rayons X dans le spectre représenté dans la partie (b)

6. Variantes de conception des compteurs proportionnels

6.1 Compteur proportionnel d'équivalent tissulaire

Un type spécialisé de compteur proportionnel connu sous le nom de compteur proportionnel équivalent tissulaire (TEPC) joue un rôle utile dans la dosimétrie des rayons gamma et des neutrons. Il est basé sur la construction d'une chambre dont les parois et le mélange de gaz de remplissage limitent la composition élémentaire du tissu biologique. Lorsqu'elles sont irradiées par des rayons gamma externes ou des neutrons, des particules chargées secondaires (principalement des électrons rapides pour les rayons gamma et des protons de recul pour les neutrons rapides) sont créées dans la paroi.

6.2 Compteurs d'avalanches à plaques parallèles

Une variante du compteur proportionnel traditionnel a attiré l'attention pour les applications dans lesquelles le besoin d'informations temporelles rapides est plus important qu'une bonne résolution en énergie.

Le détecteur d'avalanche à plaques parallèles est particulièrement utile dans les applications impliquant des particules chargées lourdes où les dommages de rayonnement associés dans les détecteurs à semi-conducteurs peuvent interdire leur utilisation. Le compteur se compose de deux électrodes plates parallèles séparées par un espace qui est normalement maintenu aussi petit que possible pour obtenir les meilleures informations de synchronisation. Alternativement, des grilles à mailles fines peuvent être utilisées comme électrodes au lieu de plaques pleines pour réduire leur épaisseur. Les électrodes sont enfermées dans un récipient dans lequel un gaz proportionnel est introduit, et un champ électrique constant est produit entre les plaques, qui, dans des conditions de pression de gaz relativement faible, peut atteindre des valeurs aussi élevées que $4 \times 10^6 \text{ V/m} \cdot \text{atm}$. Une particule chargée qui traverse l'espace entre les plaques laisse une traînée d'ions et d'électrons qui sont multipliés par le processus habituel de multiplication des gaz. Les électrons formés près de la cathode sont évidemment soumis à plus de multiplication que ceux formés près de l'anode. Des gains maximaux de l'ordre de 10^4 sont possibles.

6.3 Compteurs proportionnels sensibles à la position

Étant donné que la position de l'avalanche dans un compteur proportionnel est limitée à une petite partie de la longueur du fil d'anode, certains schémas ont été développés et qui sont capables de détecter la position d'un événement se produisant dans le tube. Dans la géométrie cylindrique commune, les électrons dérivent vers l'intérieur depuis leur lieu de formation le long des lignes de champ radiales, de sorte que la position de l'avalanche est un bon indicateur de la position axiale à laquelle la paire d'ions d'origine s'est formée. Si la trace du rayonnement incident s'étend sur une certaine distance le long de la longueur du tube, alors de nombreuses avalanches seront également réparties le long de l'anode et seule une position moyenne peut être déduite.

6.4 Compteurs proportionnels multi-fils

Dans certaines situations, il est avantageux de prévoir plus d'un fil d'anode dans un compteur proportionnel. Par exemple, des détecteurs de très grande surface peuvent être construits économiquement en plaçant une grille de fils d'anode entre deux grandes plaques plates qui servent de cathodes de chaque côté du compteur. Les électrons formés par ionisation du gaz dérivent vers l'intérieur vers le plan des fils d'anode, initialement dans un champ presque uniforme. À mesure qu'ils s'approchent, ils sont accélérés vers le fil le plus proche dans sa région de champ élevé environnant

où des avalanches se forment. Une grande impulsion induite par la polarité négative apparaît sur le fil d'anode sur lequel l'avalanche est collectée, tandis que les anodes voisines présentent des impulsions d'amplitude positive plus petites. Les signaux des préamplificateurs connectés à chaque fil peuvent être ainsi utilisés pour localiser l'événement au fil le plus proche, l'écartement des fils le plus fin étant limité à environ 1 ou 2 mm. Pour réduire le nombre de préamplificateurs, les anodes peuvent être interconnectées à l'aide de résistances pour former une chaîne de division de charge, et les signaux des préamplificateurs aux deux extrémités peuvent ensuite être traités dans le même type de schéma pour identifier le centroïde de la charge le long de la chaîne.

6.5 Compteurs proportionnels à gaz scintillant

Un détecteur hybride a été développé qui combine certaines des propriétés d'un compteur proportionnel avec celles du détecteur à scintillation (voir annexe II). Le compteur proportionnel à gaz scintillant (GPS) (également connu dans la littérature sous le nom de chambre à dérive à gaz scintillation) est basé sur la génération d'une impulsion de signal à partir des photons visibles et ultraviolets émis par des atomes ou molécules de gaz excités. Dans un scintillateur à gaz conventionnel, ces molécules de gaz excitées sont créées par l'interaction directe du rayonnement incident lorsqu'il traverse le gaz. Les durées de vie des états excités sont normalement assez courtes, de sorte que la lumière apparaît quelques nanosecondes après le temps de passage de la particule incidente.

Si un champ électrique est appliqué au gaz, les électrons libres des paires d'ions créées le long de la trajectoire des particules dériveront comme ils le font dans une chambre d'ionisation ou un compteur proportionnel. Si le champ est suffisamment fort, les collisions inélastiques entre ces électrons et les molécules de gaz neutre peuvent élever certaines des molécules à des états excités, qui peuvent alors se désexciter par l'émission d'un photon. Ce processus se produit également dans un compteur proportionnel conventionnel, mais les photons sont traités comme une nuisance car leur interaction ultérieure dans d'autres parties du compteur peut conduire à des impulsions parasites. Dans le compteur proportionnel à gaz scintillant, cependant, les photons émis par les molécules de gaz excitées sont détectés par un tube photomultiplicateur, qui produit une impulsion électrique proportionnelle au nombre de photons incidents sur le tube.

7. Détecteurs micro-modèles à gaz

À partir des années 1990, certaines des techniques impliquant la photolithographie, la gravure sélective et l'usinage au laser qui ont été initialement introduites dans l'industrie des semi-conducteurs ont commencé à être appliquées à la conception d'une nouvelle catégorie de dispositifs remplis de gaz généralement classés sous le nom de détecteurs à gaz à micro-motifs. Ils se caractérisent par des

détails fins dans leur structure et/ou leurs étages de lecture qui permettent une résolution spatiale proche de quelques micromètres, très utile dans les applications de poursuite de particules ou d'imagerie par rayonnement. La plupart ont également d'excellentes caractéristiques de synchronisation et peuvent fonctionner à des taux élevés par unité de surface. Non seulement ces propriétés présentent un grand intérêt pour la poursuite des particules à haute énergie, mais elles sont également exploitées dans un nombre croissant d'applications en dehors de la communauté de recherche en physique. De nombreuses variantes de conception ont été étudiées, et quatre configurations spécifiques qui ont atteint une large application sont données ci-dessous.

- Chambre à gaz micro-ruban ;
- Multiplicateur d'électrons de gaz ;
- Micro-mégas ;
- Chambres à plaques résistives.

Annexe I

Détecteurs à gaz (Compteurs Geiger-Mueller)

1. Introduction

Le compteur Geiger-Mueller (communément appelé compteur G-M, ou simplement tube Geiger) est l'un des plus anciens types de détecteurs de rayonnement existants, ayant été introduit par Geiger et Mueller en 1928. Cependant, la simplicité, le faible coût et la facilité de fonctionnement de ces détecteurs ont conduit à leur utilisation continue jusqu'à l'heure actuelle.

En complément des chambres d'ionisation et des compteurs proportionnels évoqués précédemment, les compteurs G-M constituent la troisième catégorie générale des détecteurs à gaz basés sur l'ionisation. En commun avec les compteurs proportionnels, ils utilisent la multiplication du gaz pour augmenter considérablement la charge représentée par les paires d'ions d'origine formées le long de la piste de rayonnement, mais d'une manière fondamentalement différente. Dans le compteur proportionnel, chaque électron d'origine conduit à une avalanche qui est fondamentalement indépendante de toutes les autres avalanches formées à partir d'autres électrons associés à l'événement ionisant d'origine. Puisque toutes les avalanches sont presque identiques, la charge collectée reste proportionnelle au nombre d'électrons d'origine.

Dans le tube G-M, des champs électriques sensiblement plus élevés augmentent l'intensité de chaque avalanche. Dans des conditions appropriées, une situation est créée dans laquelle une avalanche peut elle-même déclencher une seconde avalanche à une position différente dans le tube. A une valeur critique du champ électrique, chaque avalanche peut créer, en moyenne, au moins une avalanche de plus, et une réaction en chaîne auto-propagée en résulte. A des valeurs encore plus élevées du champ électrique, le processus devient rapidement divergent et, en principe, un nombre exponentiellement croissant d'avalanches pourrait être créé en un très peu de temps. Une fois que Geiger décharge atteint une certaine taille, cependant, les effets collectifs de toutes les avalanches individuelles entrent en jeu et finissent par mettre fin à la réaction en chaîne. Étant donné que ce point limite est toujours atteint après la création d'environ le même nombre d'avalanches, toutes les impulsions d'un tube Geiger sont de la même amplitude quel que soit le nombre de paires d'ions d'origine qui ont initié le processus. Un tube Geiger ne peut donc fonctionner que comme un simple compteur d'événements radio-induits et ne peut pas être appliqué en spectroscopie de rayonnement direct car toute information sur la quantité d'énergie déposée par le rayonnement incident est perdue. Une impulsion typique d'un tube Geiger représente une quantité inhabituellement élevée de charge collectée, environ 10^9 à 10^{10} paires d'ions se formant dans la décharge. Par conséquent, l'amplitude de l'impulsion de sortie est également importante (typiquement de l'ordre du volt). Ce signal de haut

niveau permet de simplifier considérablement l'électronique associée, éliminant souvent complètement le besoin d'amplification externe. Parce que les tubes eux-mêmes sont relativement peu coûteux, un compteur G-M est souvent le meilleur choix lorsqu'un système de comptage simple et économique est nécessaire.

Outre le manque d'informations sur l'énergie, un inconvénient majeur des compteurs G-M est leur temps mort inhabituellement long, qui dépasse largement celui de tout autre détecteur de rayonnement couramment utilisé. Ces détecteurs sont donc limités à des cadences de comptage relativement faibles, et une correction du temps mort doit être appliquée aux situations dans lesquelles le taux de comptage serait autrement considéré comme modéré (quelques centaines d'impulsions par seconde). Certains types de tubes Geiger ont également une durée de vie limitée et commenceront à tomber en panne après qu'un nombre fixe d'impulsions totales aura été enregistré.

2. La décharge Geiger

Dans une avalanche typique de Townsend créée par un unique électron d'origine, de nombreuses molécules de gaz excitées sont formées par des collisions d'électrons en plus des ions secondaires. En général pas plus de quelques nanosecondes, ces molécules excitées retournent à leur état fondamental par l'émission de photons dont la longueur d'onde peut être dans le domaine du visible ou de l'ultraviolet. Ces photons sont l'élément clé de la propagation de la réaction en chaîne qui constitue la décharge de Geiger. Un photon énergétique émis dans le remplissage d'une couche électronique interne vacante peut être réabsorbé ailleurs dans le gaz par absorption photoélectrique impliquant un électron moins étroitement lié, créant un nouvel électron libre. Alternativement, le photon peut atteindre la paroi de la cathode où il pourrait libérer un électron libre lors de l'absorption. Dans les deux cas, l'électron libre nouvellement créé migrera vers l'anode et déclenchera une autre avalanche (voir figure 01). Si le facteur multiplicateur de gaz M est relativement faible ($10^2 - 10^4$) comme dans un tube proportionnel, le nombre de molécules excitées formées dans une avalanche typique (n_0) n'est pas très grand. De plus, comme la plupart des gaz sont relativement transparents dans le visible et l'UV, la probabilité p d'absorption photoélectrique d'un photon donné est également relativement faible. Dans ces conditions, $n_0 p \ll 1$ et la situation est "sous-critique" tel que relativement peu d'avalanches se forment en plus de celles créées par les électrons libres d'origine. De nombreux gaz proportionnels contiennent également un additif pour absorber les photons préférentiellement sans création d'électrons libres, supprimant davantage la possibilité de nouvelles avalanches.

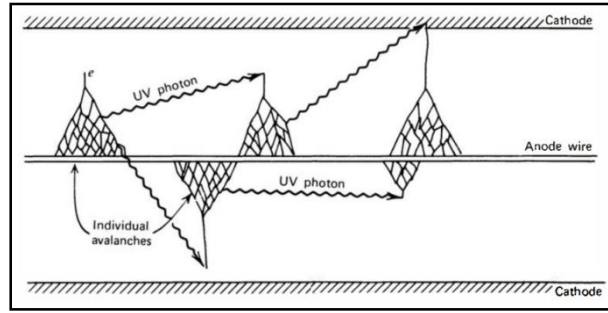


Fig. 01 Mécanisme par lequel des avalanches supplémentaires sont déclenchées par la cathode dans une décharge Geiger

Dans une décharge Geiger, cependant, la multiplication représentée par une seule avalanche est beaucoup plus élevée ($10^6 - 10^8$) et donc (n'_0) est aussi beaucoup plus grand. Désormais, les conditions de "criticité" peuvent être atteintes :

$$\boxed{n'_0 p \geq 1} \quad (01)$$

Et un nombre toujours croissant d'avalanches peut potentiellement se créer dans tout le tube.

Le temps nécessaire à la propagation de ces avalanches est relativement court. Étant donné que chaque avalanche ne commence que lorsque l'électron libre a dérivé jusqu'à quelques libres parcourus moyens du fil d'anode, le temps nécessaire pour produire toutes les paires d'ions et les molécules excitées dans une avalanche est une petite fraction de microseconde. Les photons créent préférentiellement des électrons libres proches de l'avalanche d'origine, et ces électrons doivent dériver dans la région de champ élevé près de l'anode pour générer d'autres avalanches. L'effet net est que la décharge Geiger se développe (avec certaines des caractéristiques d'une flamme chimique) dans les deux sens le long du fil avec une vitesse de propagation qui est généralement de 2 à 4 cm/ μ s, jusqu'à ce que toute la longueur du fil soit impliquée. Pour les tubes de taille habituelle, toute la charge est générée quelques microsecondes après l'événement de déclenchement et est intégré dans l'amplitude d'impulsion mesurée.

Dans un compteur proportionnel, chaque avalanche se forme à une position axiale à l'intérieur du tube qui correspond à celle de l'électron libre d'origine créé par le rayonnement incident. De plus, l'électron libre s'approche du fil d'anode à partir d'une direction radiale, et les ions secondaires sont préférentiellement formés de ce côté du fil d'anode. Cependant, dans une décharge Geiger, la propagation rapide de la réaction en chaîne conduit à de nombreuses avalanches, qui se déclenchent à des positions radiales aléatoires le long du fil. Des ions secondaires se forment donc dans toute la région multiplicatrice cylindrique qui entoure l'anode. La décharge Geiger croît donc pour envelopper tout le fil d'anode, quelle que soit la position à laquelle l'événement initiateur primaire s'est produit.

Le processus qui conduit à la fin d'une décharge Geiger a pour origine les ions positifs qui sont créés avec chaque électron dans une avalanche. La mobilité de ces ions positifs est bien inférieure à celle des électrons libres, et ils restent essentiellement immobiles pendant le temps nécessaire pour collecter tous les électrons libres de la région de multiplication. Lorsque la concentration de ces ions positifs est suffisamment élevée, leur présence commence à réduire l'amplitude du champ électrique au voisinage du fil d'anode. Puisque les ions représentent une charge d'espace positive, la région entre les ions et l'anode positive aura un champ électrique inférieur à celui prédit par l'équation (03, annexe I, détecteurs à gaz, compteur proportionnel) en l'absence de la charge d'espace. Étant donné qu'une multiplication supplémentaire du gaz nécessite le maintien d'un champ électrique au-dessus d'une certaine valeur minimale, l'accumulation de charge d'espace d'ions positifs finit par arrêter la décharge de Geiger. Une autre façon d'imaginer le mécanisme d'arrêt est de reconnaître que l'accumulation du nuage dense d'ions positifs a le même effet que l'augmentation temporaire du diamètre du fil d'anode. Le champ électrique à la surface extérieure du nuage chute alors en dessous de la valeur critique pour la poursuite de la formation d'avalanches.

Pour une tension fixe appliquée au tube, le point auquel la décharge Geiger se termine sera toujours le même, dans le sens qu'une densité donnée d'ions positifs sera nécessaire pour réduire le champ électrique en dessous de la valeur minimale requise pour une multiplication supplémentaire. Ainsi, chaque décharge Geiger se termine après avoir développé à peu près la même charge totale, quel que soit le nombre de paires d'ions d'origine créées par le rayonnement incident. Toutes les impulsions de sortie ont donc à peu près la même taille et leur amplitude ne peut fournir aucune information sur les propriétés du rayonnement incident. Lorsque la haute tension dans un tube Geiger est élevée, l'amplitude de la décharge Geiger augmente et l'amplitude de l'impulsion de sortie augmente également, comme l'indique la figure (02). Une tension plus élevée signifie un champ électrique initial plus important avant la décharge, ce qui nécessite alors une plus grande accumulation de charge d'espace avant que le champ ne soit réduit en dessous de la valeur critique. L'amplitude des impulsions augmente à peu près proportionnellement à la surtension, définie comme la différence entre la tension appliquée et la tension minimale requise pour initier une décharge Geiger.

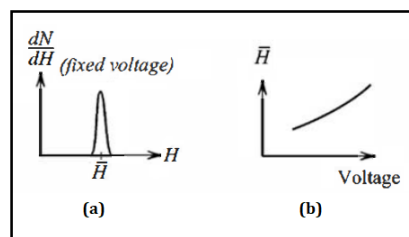


Fig. 02 (a) Spectre différentiel de hauteur d'impulsion
(b) Dépendance de l'amplitude de la décharge Geiger (amplitude de l'impulsion de sortie)
de la tension appliquée au tube Geiger

3. Gaz de remplissage

Les gaz utilisés dans un tube Geiger doivent répondre à certaines des mêmes exigences que pour les compteurs proportionnels. Étant donné que les deux types de détecteurs sont basés sur la multiplication des gaz, même des traces de gaz qui forment des ions négatifs (tels que l'oxygène) doivent être évitées dans les deux cas. Les gaz nobles sont largement utilisés comme composant principal du gaz de remplissage dans les tubes G-M, l'hélium et l'argon étant les choix les plus populaires. Un deuxième composant est normalement ajouté à la plupart des gaz Geiger dans le but d'extinction.

L'énergie moyenne atteinte par un électron libre entre les collisions avec des molécules de gaz dépend du rapport de l'amplitude du champ électrique à la pression du gaz, $\mathcal{E}/p$.

Étant donné que le nombre de molécules excitées n devrait augmenter également avec ce paramètre, le développement d'une décharge Geiger complète nécessite une certaine valeur minimale de $\mathcal{E}/p$, qui dépend du mélange du gaz spécifique utilisé dans le tube G-M. Il est courant de concevoir des tubes scellés avec un gaz de remplissage maintenu à une pression inférieure à la pression atmosphérique afin de minimiser la valeur de $\mathcal{E}$ nécessaire pour atteindre le rapport $\mathcal{E}/p$ voulu. Pour les gaz normaux à des pressions de plusieurs dixièmes d'un « atmosphère », des tensions de fonctionnement d'environ 500 - 2000 V sont nécessaires pour atteindre l'intensité de champ nécessaire en utilisant des fils d'anode d'environ 0,1 mm de diamètre. Les tubes Geiger peuvent être également conçus pour fonctionner à la pression atmosphérique. Dans de tels cas, la conception mécanique peut être simplifiée (en particulier pour les grands tubes) car aucun différentiel de pression n'a besoin d'être supporté à travers les parois ou la fenêtre du tube. Tout comme dans les compteurs proportionnels, soit le gaz peut être scellé de manière permanente dans le tube, soit des dispositions peuvent être prises pour faire circuler le gaz à travers le volume du compteur pour reconstituer le gaz en continu et éliminer toutes les impuretés.

4. Extinction

Si un tube Geiger est rempli d'un seul gaz tel que l'argon, alors tous les ions positifs formés sont des ions de cette même espèce gazeuse. Une fois la décharge Geiger primaire terminée, les ions positifs s'éloignent lentement du fil d'anode et arrivent finalement à la cathode ou à la paroi extérieure du compteur. Ils sont neutralisés en se combinant avec un électron de la surface de la cathode. Dans ce processus, une quantité d'énergie égale à l'énergie d'ionisation du gaz moins l'énergie nécessaire pour extraire l'électron de la surface de la cathode (la fonction de travail) est libérée. Si cette énergie libérée dépasse également le travail de sortie de la cathode, il est énergétiquement possible qu'un autre électron libre émerge de la surface de la cathode. Ce sera le cas si l'énergie d'ionisation du gaz dépasse le double de la valeur du travail de sortie. La probabilité est toujours faible qu'un ion donné libère un

électron lors de sa neutralisation, mais si le nombre total d'ions est suffisamment grand, il y aura probablement au moins un tel électron libre généré. Cet électron dérivera alors vers l'anode et déclenchera une autre avalanche, conduisant à une seconde décharge complète de Geiger. Le cycle entier sera maintenant répété et, dans ces circonstances, le compteur G-M, une fois initialement déclenché, produirait une sortie continue de multiples d'impulsions.

Le problème des impulsions multiples est potentiellement beaucoup plus grave dans les tubes Geiger que dans les compteurs proportionnels. Dans ce dernier cas, le nombre d'ions positifs par événement est suffisamment petit pour que seule une impulsion parasite occasionnelle puisse en résulter. L'amplitude de l'impulsion parasite ne correspondra qu'à une seule avalanche et sera généralement considérablement plus faible que celle produite par l'événement initial. Avec un tube Geiger, le nombre accru d'ions positifs atteignant la cathode améliore considérablement la probabilité d'émission d'électrons libres. De plus, la décharge secondaire, bien que provoquée par un seul électron, atteint une amplitude égale à celle de la décharge primaire.

Pour ces raisons, des précautions particulières doivent être prises dans les compteurs Geiger pour éviter la possibilité d'impulsions multiples excessives. L'extinction externe consiste en une méthode pour réduire la haute tension appliquée au tube, pendant un temps fixe après chaque impulsion, à une valeur trop faible pour supporter une multiplication supplémentaire du gaz. Ensuite, des avalanches secondaires ne peuvent pas se former, et même si un électron libre est libéré à la cathode, il ne peut pas provoquer une autre décharge du tube Geiger.

La tension doit être supprimée pendant un temps supérieur au temps de transit de l'ion positif de son lieu de formation à la cathode (généralement quelques centaines de microsecondes), auquel s'ajoute le temps de transit de l'électron libre (généralement de l'ordre de la microseconde).

Une méthode d'extinction externe consiste simplement à choisir R sur la figure (03) comme étant une valeur suffisamment grande (généralement 10^8 ohms) pour que la constante de temps du circuit de collecte de charge soit de l'ordre de la milliseconde. Cette méthode d'extinction par « Résistance externe » a l'inconvénient de nécessiter plusieurs millisecondes pour que l'anode revienne près de sa tension normale, et donc des décharges Geiger complètes pour chaque événement ne sont produites qu'à des taux de comptage très faibles.

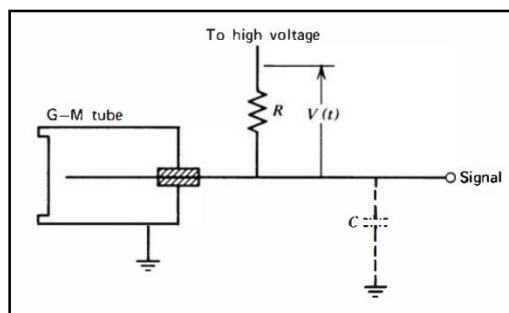


Fig. 03 Circuit de comptage équivalent pour un tube G-M. Le produit RC (généralement uniquement la capacité inhérente du tube et de l'électronique) détermine la constante de temps de la restauration de la haute tension suite à une décharge Geiger

Il est beaucoup plus courant d'empêcher la possibilité d'impulsions multiples par une extinction interne, ce qui est accompli en ajoutant un deuxième composant appelé gaz d'extinction au gaz de remplissage primaire. Il est choisi pour avoir un potentiel d'ionisation plus faible et une structure moléculaire que le composant gazeux primaire et est présent à une concentration typique de 5 à 10 %. Bien qu'il porte le même nom que le composant ajouté aux gaz de compteur proportionnels pour absorber les photons UV, le gaz d'extinction dans un compteur Geiger a une fonction différente. Il empêche les impulsions multiples grâce au mécanisme des collisions de transfert de charge. Les ions positifs formés par le rayonnement incident sont pour la plupart du composant primaire, et ils font ensuite de nombreuses collisions avec des molécules de gaz neutre lorsqu'ils dérivent vers la cathode. Certaines de ces collisions se produiront avec des molécules du gaz d'extinction et, en raison de la différence des énergies d'ionisation, il y aura une tendance à transférer la charge positive à la molécule du gaz d'extinction. L'ion positif d'origine est ainsi neutralisé par transfert d'un électron et un ion positif du gaz d'extinction commence à dériver à sa place. Si la concentration du gaz d'extinction est suffisamment élevée, ces collisions de transfert de charge garantissent que tous les ions qui arriveront éventuellement à la cathode seront ceux du gaz d'extinction. Lorsqu'elles sont neutralisées, l'excès d'énergie peut maintenant servir à la dissociation des molécules plus complexes qu'à la libération d'un électron libre de la surface de la cathode. Avec le bon choix de gaz d'extinction, la probabilité de dissociation peut être beaucoup plus grande que celle de l'émission d'électrons, et donc aucune avalanche supplémentaire n'est formée à l'intérieur du tube.

De nombreuses molécules organiques possèdent les caractéristiques appropriées pour servir de gaz d'extinction. Parmi lesquelles, l'alcool éthylique et le formiate d'éthyle se sont être révélés les plus populaires et sont largement utilisés dans de nombreux tubes G-M commerciaux. Du fait que le gaz d'extinction est dissocié pendant l'exercice de sa fonction, il est progressivement consommé au cours de la durée de vie du tube. Un tube trempé organique peut généralement avoir une limite

d'environ 10^9 comptes dans sa durée de vie utile avant que des impulsions multiples puissent ne plus être empêchées en raison de l'épuisement du gaz d'extinction.

Pour éviter le problème de la durée de vie limitée, certains tubes utilisent des halogènes (chlore ou brome) comme gaz d'extinction. Bien que les molécules d'halogène se dissocient également lors de l'exécution de leur fonction d'extinction, elles peuvent être reconstituées ultérieurement par recombinaison spontanée. En principe, les tubes à extinction d'halogènes ont donc une durée de vie infinie et sont privilégiés dans les situations où l'application implique une utilisation prolongée à des cadences de comptage élevées. Cependant, il existe d'autres mécanismes que la consommation du gaz d'extinction qui limitent la durée de vie des tubes Geiger. Deux des plus importants semblent être la contamination du gaz par les produits de réaction produite dans la décharge et les changements sur la surface de l'anode provoqués par le dépôt de produits de réaction polymérisés.

5. Comportement temporel

5.1 Profil d'impulsions

Dans les compteurs proportionnels le développement de l'impulsion d'une seule avalanche dérive presque entièrement du mouvement des ions positifs loin du fil d'anode. Une composante initiale rapide (avec un temps de montée, inférieur à une microseconde) résulte du mouvement des ions à travers la région de champ élevé près du fil d'anode. Une augmentation beaucoup plus lente s'ensuit qui correspond à une dérive plus lente des ions à travers la région de champ de faible intensité qui existe dans la majeure partie du volume du tube. Dans une décharge Geiger, l'impulsion correspond à l'effet cumulatif de nombreuses avalanches qui se sont formées sur toute la longueur du fil d'anode. Le temps nécessaire aux électrons secondaires pour atteindre la région de multiplication pour déclencher ces multiples avalanches sera variable, mais les différences seront limitées par le mouvement relativement rapide des électrons. La montée initiale de l'impulsion sera un peu plus lente que celle d'une seule avalanche, mais se développera toujours sur un temps généralement inférieur à quelques microsecondes. La situation est également compliquée par les distorsions de champ électrique qui se produisent à partir de l'accumulation de charge d'espace près de la fin de la décharge.

L'effet net est que l'impulsion de sortie d'un tube G-M présente toujours une augmentation rapide de l'ordre d'une microseconde ou moins, suivie d'une augmentation beaucoup plus lente en raison de la lente dérive des ions. Si la constante de temps du circuit de collecte était infinie, cela prendrait tout le temps de collecte des ions pour que l'impulsion de sortie atteigne son maximum théorique. En tenant compte du taux de comptage maximal, on choisit souvent des constantes de temps bien inférieures à $100\ \mu\text{s}$, ce qui élimine largement la partie à montée lente de l'impulsion et ne laisse que le front montant rapide, comme illustré à la figure (04). Même si une fraction significative de l'amplitude de l'impulsion de sortie potentielle peut être perdue, la grande quantité de charge générée dans la décharge Geiger produit une impulsion si grande qu'une certaine perte d'amplitude peut facilement être tolérée. De plus, étant donné que toutes les décharges Geiger sont approximativement uniformes en taille et en profil temporel, toutes les impulsions seront atténuées par le

même taux dans le processus de mise en forme et les impulsions de sortie resteront presque de la même amplitude.

Comme dans les compteurs proportionnels, il existe un délai entre le moment de la formation de la première paire d'ions dans le gaz et le déclenchement de la première avalanche. Ce retard correspond au temps de dérive de l'électron de son point d'origine à la région de multiplication et peut atteindre une microseconde ou plus dans les grands tubes. Lorsque des tubes G-M sont utilisés pour le chronométrage, il faut donc s'attendre à une incertitude temporelle de cette ampleur si des interactions de rayonnement se produisent de manière aléatoire dans tout le volume du tube.

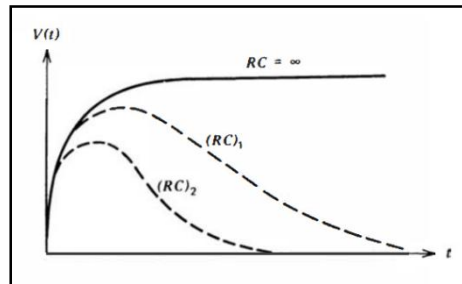


Fig. 04 Forme de l'impulsion de sortie du tube G-M pour différentes constantes de temps RC du circuit de comptage

$(RC)_2 < (RC)_1 < \text{constante de temps infinie}$. Le signal de la tension $V(t)$ est supposé être mesuré comme indiqué sur la figure (03), donnant une polarité positive

5.2 Temps mort

L'accumulation de la charge d'espace ionique positive qui arrête la décharge Geiger garantit également qu'un temps considérable doit s'écouler avant qu'une seconde décharge Geiger puisse être générée dans le tube. Au fur et à mesure que les ions positifs dérivent radialement vers l'extérieur, la charge d'espace devient plus diffusante et le champ électrique dans la région de multiplication commence à revenir à sa valeur d'origine.

Une fois que les ions positifs ont parcouru une partie de la distance, le champ aura suffisamment récupéré son intensité pour permettre une autre décharge Geiger. Cependant, si le champ n'a pas été entièrement restauré, la décharge sera moins intense que l'original car moins d'ions positifs seront nécessaires pour arrêter la décharge en réduisant à nouveau le champ électrique en dessous du point critique. Les premières impulsions qui apparaissent sont donc réduites en amplitude par rapport à l'originale et peuvent ou non être enregistrées par le système de comptage, selon sa sensibilité. Lorsque les ions positifs ont finalement dérivé jusqu'à la cathode, le champ électrique est entièrement restauré et un autre événement ionisant déclenchera une seconde décharge Geiger de pleine amplitude. Ce comportement est illustré sur la figure (05), qui est similaire à une trace d'oscilloscope qui peut être observée expérimentalement à partir de tubes Geiger à des taux de comptage élevés.

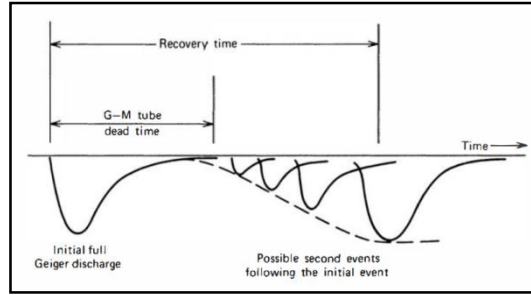


Fig. 05 Illustration du temps mort d'un tube G-M

Les impulsions de polarité négative observées de manière conventionnelle à partir du détecteur sont affichées. Immédiatement après la décharge de Geiger, le champ électrique a été réduit en dessous du point critique par la charge d'espace positive. Si un autre événement ionisant se produit dans ces conditions, une deuxième impulsion ne sera pas observée car la multiplication des gaz est empêchée. Pendant ce temps, le tube est donc "mort" et toutes les interactions de rayonnement qui se produisent dans le tube pendant ce temps seront perdues. Techniquement, le temps mort du tube Geiger est défini comme la période entre l'impulsion initiale et le moment où une deuxième décharge Geiger, quelle que soit sa taille, peut être développée.

Dans la plupart des tubes Geiger, ce temps est de l'ordre de 50 -100 μ s. Dans tout système de comptage pratique, une certaine amplitude d'impulsion finie doit être atteinte avant que la seconde impulsion ne soit enregistrée, et le temps écoulé nécessaire pour développer une seconde décharge qui dépasse cette amplitude est parfois appelé le temps de résolution du système. En pratique, ces deux termes sont souvent utilisés de manière interchangeable et le terme temps mort peut être également utilisé pour décrire le comportement combiné du système détecteur-comptage. Le temps de récupération est l'intervalle de temps nécessaire pour que le tube revienne à son état d'origine et devient capable de produire une seconde impulsion de pleine amplitude.

Le comportement du temps mort du détecteur G-M est spécifique, car le tube G-M est un exemple rare de type de détecteur de rayonnement qui présente un long temps mort inhérent résultant des mécanismes qui se produisent dans le détecteur lui-même. Dans la plupart des autres cas, les processus internes au détecteur n'empêchent pas la production d'un signal de sortie du détecteur même pour deux événements qui peuvent se produire avec un espacement temporel infiniment petit. Par exemple, si deux particules interagissent étroitement espacées dans le temps dans une chambre ionique de type impulsionnel, les charges des deux pistes dériveront indépendamment et le signal induit observé sera une superposition des signaux attendus séparément de chacun. Si ces deux composants chevauchent suffisamment, alors une seule impulsion peut être enregistrée et une perte se produira qui est aussi généralement attribuée au temps mort du système. Les détails de ces pertes

dépendent maintenant de la façon dont les impulsions du détecteur sont formées, traitées ou comptées, de sorte que le comportement de temps mort de la plupart des systèmes de mesure dépend de manière critique de l'électronique de traitement des impulsions ainsi que des caractéristiques temporelles des signaux du détecteur. Les mêmes commentaires généraux s'appliquent aux détecteurs à scintillation et à semi-conducteurs traités plus loin dans ces annexes.

La plupart des applications de ces détecteurs impliquent la spectroscopie plutôt que la simple opération de comptage à laquelle les tubes G-M sont limités.

6. Le plateau de comptage de Geiger

Le tube Geiger fonctionnant comme un simple compteur, son application ne nécessite que l'établissement des conditions de fonctionnement dans lesquelles chaque impulsion est enregistrée par le système de comptage. En pratique, ce point de fonctionnement est normalement choisi en enregistrant une courbe plateau du système dans des conditions dans lesquelles une source de rayonnement génère des événements à vitesse constante à l'intérieur du tube.

Le taux de comptage est enregistré au fur et à mesure que la haute tension appliquée au tube croît à partir d'une valeur initialement basse. Les résultats sont particulièrement simples pour le cas du tube Geiger dans lequel il s'agit d'impulsions d'amplitudes centrées autour d'une seule valeur. La distribution différentielle d'hauteur d'impulsion a alors l'allure simple donnée dans la figure (06 a) pour une valeur fixe de la tension appliquée au tube. Si la tension augmente, l'amplitude d'impulsion moyenne augmente également et le pic de la distribution de hauteur d'impulsion différentielle se décale vers la droite. Les distributions sont indiquées pour deux valeurs hypothétiques de tension appliquée dans la figure.

Si l'amplitude minimale requise par le système de comptage est H_d comme indiqué sur la figure, alors aucune impulsion ne doit être comptée lorsque la tension appliquée est de 1000 V, mais toutes les impulsions doivent être comptées à 1200 V. Le taux de comptage correspondant passe donc brusquement de zéro à la valeur maximale entre ces deux valeurs de la tension. Augmenter davantage la tension ne fait que déplacer le pic plus vers la droite dans la distribution différentielle, mais n'augmente pas le nombre d'événements enregistrés. Par conséquent, la courbe de comptage montre un plateau plat à des tensions supérieures à 1200 V, et en principe tout choix de tension de fonctionnement au-dessus de cette valeur assurerait un fonctionnement stable. La tension minimale à laquelle les impulsions sont enregistrées pour la première fois par le système de comptage est souvent appelée tension de démarrage, tandis que la transition entre la montée rapide de la courbe et le plateau est son « genou ».

Si la tension est suffisamment élevée, le plateau se termine brusquement en raison de l'apparition de mécanismes de décharge continue à l'intérieur du tube. Il peut s'agir de décharges corona déclenchées

par des irrégularités nettes sur le fil d'anode, ou de multiples impulsions causées par une défaillance du mécanisme d'extinction. Le processus de décharge continue est potentiellement nocif s'il est maintenu pendant une durée quelconque, et il faut donc immédiatement diminuer la tension appliquée lorsque la fin du plateau est observée. Dans l'intérêt d'une longue durée de vie, on choisit normalement un point de fonctionnement qui n'est que suffisamment haut sur le plateau pour garantir que la région plate a été atteinte. Les décharges Geiger sont ainsi maintenues proches de la taille minimale compatible avec une bonne stabilité de comptage et le gaz d'extinction est consommé au débit minimal.

Dans les cas réels, le plateau de comptage montre toujours une pente finie, comme le montre la figure (06 b). Tout effet qui ajoute une queue de faible amplitude à la distribution différentielle de hauteur d'impulsion peut être une cause contributive de la pente. Par exemple, certaines régions proches des extrémités du tube peuvent avoir une intensité de champ électrique inférieure à la normale et les décharges provenant de ces régions peuvent être plus petites que la normale. En outre, toutes les impulsions qui se produisent pendant le temps de récupération seront également anormalement petites. Afin d'exclure ce dernier effet lors de la mesure des caractéristiques de plateau, le taux de comptage ne doit pas dépasser quelques pourcents de l'inverse du temps de récupération. Pour les tubes G-M typiques, cette restriction signifie que la mesure du plateau doit être effectuée à une vitesse inférieure à quelques centaines de coups par seconde.

Une autre cause de la pente dans le plateau de nombreux tubes G - M est la défaillance occasionnelle du mécanisme d'extinction, qui peut conduire à une impulsion satellite ou parasite en plus de la décharge Geiger primaire. La probabilité d'un tel événement est normalement faible mais augmentera au fur et à mesure que le nombre d'ions positifs qui atteignent finalement la cathode augmente également. Le volume actif du tube peut également augmenter légèrement avec la tension car le champ électrique très faible près des coins du tube augmente et la recombinaison de paires d'ions est ainsi empêchée.

La courbe plateau peut également présenter un léger effet d'hystérésis, est que la courbe de comptage enregistrée en augmentant la tension peut différer quelque peu de celle obtenue en diminuant la tension. Cette hystérésis survient parce que les charges électriques sur les isolants sont lentes à s'équilibrer et peuvent influencer la configuration du champ électrique à l'intérieur du tube.

Les tubes Geiger sont souvent évalués en fonction de la pente indiquée par la partie linéaire du plateau.

Les tubes à extinction organiques ont tendance à présenter le plateau le plus plat, avec des pentes de l'ordre de 2 à 3 % par changement de 100 V de la tension appliquée. Les tubes halogènes à

extinction ont un plateau de pente typique plus important, mais leur durée de vie utile plus longue peut parfois compenser cet inconvénient. Les tubes halogènes peuvent également fonctionner à des faibles valeurs des tensions appliquées. Par exemple, les compteurs argon-alcool nécessitent généralement une tension de démarrage d'au moins 1000 V, mais les tubes remplis de néon à extinction avec du brome peuvent afficher une tension de démarrage aussi basse que 200-300 V. Les mécanismes qui conduisent à cette tension de démarrage basse dépendent fortement de la pression partielle de l'halogène et de la valeur $\mathcal{E}/P$ à l'intérieur du tube. L'interaction des atomes de brome avec les atomes métastables et excités de néon jouent apparemment un rôle important dans la propagation de la décharge et dans la détermination des propriétés du tube résultant.

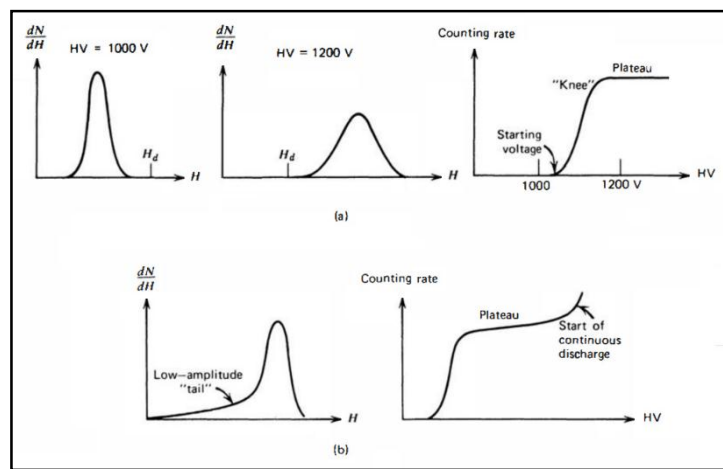


Fig. 06 (a) Etablissement du plateau de comptage pour un tube G-M. Comme la haute tension varie dans cet exemple de 1000 à 1200 V, les impulsions de sortie passent d'une chute en dessous du seuil de compteur H_d à une situation dans laquelle toutes les impulsions sont supérieures à H_d .
(b) La queue de faible amplitude sur le spectre de hauteur d'impulsion à gauche provoque une pente finie du plateau dans la courbe de comptage.

7. Caractéristiques de conception

Les tubes Geiger partagent certaines caractéristiques générales de conception avec les compteurs proportionnels, mais ils sont moins exigeants. Puisque la multiplication des gaz et la formation d'avalanches sont à nouveau importantes, l'utilisation d'un fil d'anode fin est presque universelle afin de produire les hautes valeurs locales du champ électrique nécessaires pour créer une région multiplicatrice. Dans le tube Geiger, moins d'attention à l'uniformité du fil est nécessaire car la décharge Geiger se propage sur toute la longueur de l'anode, et les non-uniformités sont donc moyennées. Comme l'amplitude de l'impulsion ne contient aucune information quantitative, l'uniformité de la réponse est moins importante et, par conséquent, les tubes de champ utilisés dans les compteurs proportionnels pour façonner le champ électrique se trouvent rarement dans les tubes

Geiger. Certains modèles peuvent être utilisés de manière interchangeable comme des tubes proportionnels ou Geiger, bien qu'un changement du gaz de remplissage soit normalement recommandé pour des performances optimales.

Une conception typique des tubes Geiger est le type de fenêtre d'extrémité illustré à la figure (07). Le fil d'anode est supporté à une seule extrémité et est situé le long de l'axe d'une cathode cylindrique en métal ou en verre avec un revêtement intérieur métallisé. Le rayonnement pénètre dans le tube par la fenêtre d'entrée, qui peut être en mica ou en un autre matériau capable de conserver sa résistance dans des sections minces. Étant donné que la plupart des tubes Geiger fonctionnent en dessous de la pression atmosphérique, la fenêtre doit supporter une pression différentielle substantielle. La fenêtre doit être aussi fine que possible lors du comptage des particules à courte portée, telles que les alphas, mais peut être rendue plus robuste pour les applications impliquant des particules bêta ou des rayons gamma.

Les tubes Geiger peuvent être également fabriqués dans de nombreuses autres formes et configurations. Une simple anode à boucle de fil insérée dans un volume arbitraire entouré d'une cathode conductrice fonctionnera normalement comme un compteur Geiger, mais il faut prendre soin d'éviter les régions à faible champ dans les coins où des événements peuvent être perdus. Les compteurs Geiger peuvent également prendre la forme de « compteurs à aiguilles », dans lesquels l'anode consiste en une aiguille pointue soutenue dans un gaz enfermé par la cathode. Le champ au voisinage du point varie en $1/r^2$, plutôt qu'en $1/r$ comme dans la géométrie cylindrique courante. Les compteurs avec de petits volumes actifs peuvent être facilement fabriqués en utilisant la géométrie de l'aiguille.

Pour les applications impliquant des rayonnements très doux comme les particules chargées lourdes de faible énergie ou les particules « bêta » molles, il peut être préférable d'introduire la source directement dans le volume de comptage. Les compteurs Geiger à débit continu, dont la conception est similaire aux compteurs proportionnels à débit de gaz, sont souvent utilisés pour ces raisons. En faisant circuler en continu le gaz du compteur à travers la chambre, la pureté du gaz est maintenue malgré de petites quantités d'air introduite lors du changement d'échantillons.

Des rendements de comptage très élevés pour certains rayonnements doux peuvent être obtenus en introduisant la source sous forme d'un gaz directement mélangé au gaz du compteur. Par exemple, de faibles concentrations de ^{14}C peuvent être détectées par conversion du carbone en CO_2 , qui peut être ensuite mélangé avec un gaz de remplissage argon-alcool classique. Les propriétés d'un gaz Geiger sont conservées à condition que la concentration de CO_2 ne soit pas trop élevée.

Les particules bêta molles du tritium peuvent également être efficacement comptées à l'aide d'un gaz de remplissage constitué d'hydrogène dans lequel une petite quantité de vapeur d'eau a été ajoutée.

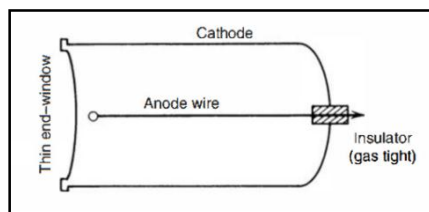


Fig. 07 Coupe transversale d'un tube Geiger à fenêtre d'extrémité

Comme aucun préamplificateur n'est normalement requis, un système de comptage impliquant des tubes Geiger peut être aussi simple que celui illustré à la figure (08). La résistance série R entre l'alimentation haute tension et l'anode du tube est la résistance de charge à travers laquelle la tension du signal est développée. La combinaison parallèle de R avec la capacité du tube et du câblage associé (indiqué par C_s) détermine la constante de temps du circuit de collecte de charge. Comme vu dans la section concernant la réponse temporelle, cette constante de temps est normalement choisie pour être de quelques microsecondes afin que seules les composantes à montée rapide de l'impulsion soient conservées. Le condensateur de couplage C_c est également représenté, ce qui est nécessaire pour bloquer la haute tension du tube mais transmettant l'impulsion du signal aux circuits suivants. Pour réaliser cette fonction sans atténuer l'amplitude de l'impulsion, sa valeur doit être suffisamment grande pour que $R.C_c$ soit grand devant la durée de l'impulsion.

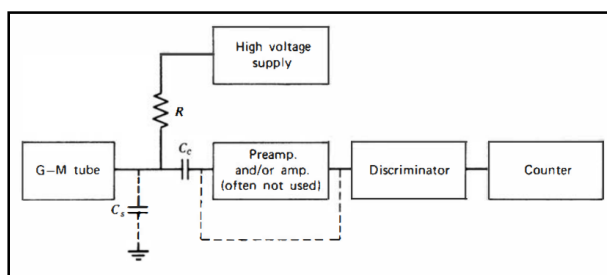


Fig. 08 Schéma fonctionnel de l'électronique de comptage associée à un tube G-M

8. Efficacité de comptage

8.1 Particules chargées

Étant donné qu'une seule paire d'ions formée dans le gaz de remplissage du tube Geiger peut déclencher une décharge Geiger complète, l'efficacité de comptage pour toute particule chargée qui pénètre dans le volume actif du tube est essentiellement de 100 %. Dans la plupart des situations pratiques, l'efficacité effective de comptage est donc déterminée par la probabilité que le rayonnement incident pénètre dans la fenêtre du tube sans absorption ni rétrodiffusion. Avec les particules alpha, l'absorption à l'intérieur de la fenêtre est la principale préoccupation à tenir en compte et des fenêtres d'une épaisseur aussi petite que $1,5 \text{ mg/cm}^2$ sont commercialisées. Des fenêtres plus épaisses peuvent généralement être utilisées pour compter les particules « bêta », mais certaines seront rétrodiffusées par la fenêtre sans pénétrer si l'épaisseur de la fenêtre est une fraction significative de la portée des électrons.

8.2 Neutrons

Il y a plusieurs raisons pour lesquelles les tubes Geiger sont rarement utilisés pour compter les neutrons. Pour les neutrons thermiques, les gaz Geiger classiques ont de faibles sections efficaces de capture et ont donc une efficacité de comptage très faible. Les gaz avec une section efficace de capture élevée (tels que ^3He) peuvent être utilisés, mais alors le détecteur fonctionne de manière beaucoup plus sensible dans la région proportionnelle que dans la région de Geiger. Dans le compteur proportionnel, les événements induits par les neutrons ont une amplitude beaucoup plus grande que les impulsions générées par le fond gamma et sont donc faciles à distinguer. Dans la région de Geiger, toutes les impulsions ont la même amplitude et la discrimination des rayons gamma n'est pas possible.

Les neutrons rapides peuvent produire des noyaux de recul dans un gaz Geiger qui génèrent des paires d'ions et une décharge ultérieure. Par conséquent, les tubes Geiger répondront, dans une certaine mesure, à un flux de neutrons rapides, en particulier ceux qui sont remplis d'hélium. Cependant, les détecteurs de neutrons rapides remplis de gaz fonctionnent normalement comme des compteurs proportionnels plutôt que comme des tubes Geiger pour tirer parti des informations spectroscopiques fournies uniquement dans la région proportionnelle.

8.3 Rayons gamma

Dans tout compteur rempli de gaz, la réponse aux rayons gamma d'énergie normale se produit au moyen d'interactions de rayons gamma dans la paroi solide du compteur. Si l'interaction a lieu suffisamment près de la surface de la paroi interne pour que l'électron secondaire créé dans l'interaction puisse atteindre le gaz et créer des ions, une impulsion en résultera. Puisqu'une seule

paire d'ions est nécessaire, l'électron secondaire n'a qu'à peine besoin d'émerger de la paroi près de l'extrémité de sa piste pour générer une impulsion à partir d'un tube Geiger.

L'efficacité du comptage des rayons gamma dépend donc de deux facteurs distincts :

1. La probabilité que le rayon gamma incident interagisse dans la paroi et produise un électron secondaire ;
2. La probabilité que l'électron secondaire atteigne le gaz de remplissage avant la fin de sa piste.

Comme le montre la figure (09), seule la couche la plus loin de la paroi près du gaz peut apporter des électrons secondaires. Cette région a une épaisseur égale à la portée maximale des électrons secondaires qui se forment, soit typiquement un millimètre ou deux. Rendre la paroi plus épaisse que cette valeur ne contribue en rien à l'efficacité car les électrons formés dans les régions de la paroi plus éloignées du gaz n'ont aucune chance d'atteindre le gaz avant d'être arrêtés.

La probabilité d'interaction des rayons gamma dans cette couche critique augmente généralement avec le numéro atomique du matériau de la paroi. Par conséquent, l'efficacité des rayons gamma des tubes Geiger est la plus élevée pour les tubes construits avec une paroi cathodique en matériau à Z élevé. Les tubes G-M avec des cathodes au bismuth ($Z = 83$) ont été largement utilisés pour la détection des rayons gamma. La probabilité d'interaction au sein de la couche sensible reste faible même pour les matériaux à Z élevé, et les efficacités typiques de comptage des rayons gamma sont rarement supérieures à quelques pour cent.

Si l'énergie des photons est suffisamment faible, l'interaction directe des rayons gamma dans le gaz de remplissage du tube peut ne plus être négligeable. Par conséquent, l'efficacité de comptage des rayons gamma et des rayons X de faible énergie peut être améliorée en utilisant des gaz de numéro atomique élevé à une pression aussi élevée que possible.

Le xénon et le krypton sont des choix courants dans ces applications, et des efficacités de comptage proches de 100 % peuvent être atteintes pour des énergies de photons inférieures à environ 10 ke V.

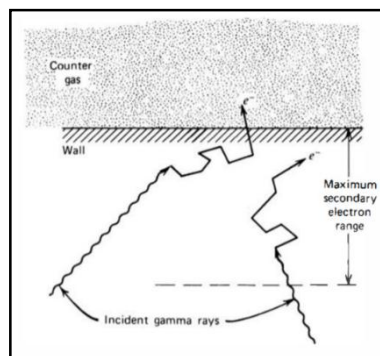


Fig. 09 Mécanisme principal par lequel les compteurs remplis de gaz sont sensibles aux rayons gamma implique la création d'électrons secondaires dans la paroi du compteur. Seules les interactions qui se produisent dans une plage d'électrons de la surface de la paroi peuvent entraîner une impulsion.

9. Méthode de « time-to-first-count »

En raison du long temps mort intrinsèque des tubes Geiger, ils sont normalement limités aux applications dans lesquelles le taux de comptage est relativement faible. Les corrections pour les pertes de temps mort deviennent significatives même à des taux de quelques centaines de comptes par seconde. À des taux supérieurs à plusieurs milliers par seconde, ces corrections deviennent généralement si importantes que le véritable taux de comptage devient très incertain et extrêmement sensible à une faible variation du temps mort ou du modèle utilisé pour représenter ces pertes.

Afin d'étendre la plage de comptage effective des tubes Geiger au-delà de cette limite, un mode de fonctionnement innovant a été introduit, connu sous le nom de méthode du premier comptage. Comme le montre la figure (10), la méthode est basée sur la commutation la haute tension appliquée au tube Geiger entre deux valeurs : la tension normale de fonctionnement, et une tension inférieure qui est réglée pour être inférieure à la tension minimale nécessaire pour développer des avalanches dans le gaz de remplissage. Après une décharge Geiger, la tension chute brusquement à la valeur inférieure et y est maintenue pendant un "temps d'attente" fixe (généralement 1 à 2 millisecondes). Ce temps d'attente est choisi supérieur au temps de rétablissement du tube à la tension réduite appliquée. Pendant le temps d'attente, les ions positifs générés dans la décharge Geiger précédente sont balayés vers la cathode et le tube est restauré dans son état normal ou "sous tension". A la fin de cette période, la tension est à nouveau brusquement commutée vers la valeur supérieure de fonctionnement. Une fois cette commutation effectuée, le tube est par définition dans un état « sous tension » puisque tous les ions ont été éliminés. A ce moment, une horloge démarre, et le temps écoulé est ensuite enregistré jusqu'à ce que la prochaine décharge Geiger se produise. La tension appliquée est à nouveau diminuée pendant le temps d'attente fixé et le cycle est répété plusieurs fois.

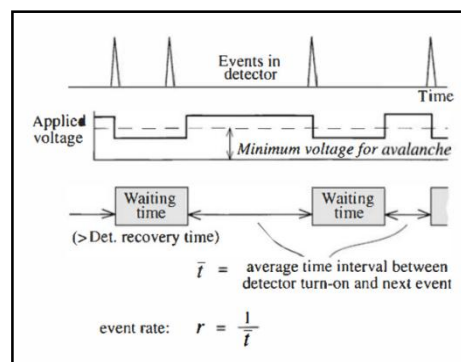


Fig. 10 Illustration de la méthode « time-to-first-count » pour déterminer le taux d'événements à partir d'un tube G-M

Grâce à ce processus, une série de mesures sont effectuées sur le temps entre la fin du temps d'attente et la prochaine occurrence d'un événement réel dans le tube. Lorsque ces intervalles de temps

sont moyennés, on a alors une mesure du temps moyen entre un point arbitraire sur l'échelle de temps et l'événement vrai suivant. Il a été montré que de tels intervalles ont une longueur moyenne donnée par $1/r$, où r est le vrai taux d'événements. Ainsi, ce taux peut être déterminé simplement en prenant l'inverse de l'intervalle de temps moyen mesuré.

A noter que cette technique est totalement exempte des effets du temps mort du détecteur. Chaque période de mesure du temps commence avec le détecteur dans un état actif et se termine lorsque le prochain événement réel se produit. La précision de la mesure n'est limitée que par la brusque commutation de la haute tension qui doit être utilisée pour définir le temps zéro et l'incertitude sur le moment du début de la décharge de Geiger. Ces incertitudes peuvent être aussi courtes qu'un dixième de microseconde, de sorte que des intervalles de temps moyens de quelques microsecondes ou plus peuvent être mesurés avec une précision raisonnable.

Cet intervalle de temps moyen correspond à l'extension du taux de comptage mesuré à plus de 10^5 coups par seconde, soit plus de deux décades de plus que permis par le fonctionnement classique du tube.

Il y aura également des incertitudes statistiques associées à ce type de mesure. En général, le nombre d'intervalles de temps mesurés jouera le même rôle que le nombre de coups enregistrés dans une mesure de comptage classique sur une période de temps fixe. Dans ce cas, un nombre de comptages égal au nombre d'intervalles mesurés est enregistré sur une durée de vie totale égale au produit du nombre d'intervalles multiplié par leur durée moyenne.

Pour obtenir une précision statistique de 1 %, il faut donc mesurer au moins 10^4 intervalles de temps. Cependant, étant donné que la méthode est le plus souvent appliquée à des circonstances de taux de comptage élevé, l'enregistrement d'un grand nombre d'intervalles sera assuré dans un temps total de mesure raisonnablement court. Dans la limite des taux vrais très élevés, le taux d'intervalles mesurés se rapproche de l'inverse du temps d'attente fixe choisi. Si cette valeur est de 2 ms, alors le taux d'acquisition de données s'approche de 500/s.

10. Compteur de balayage G-M

Un type courant de compteur utilisé dans le balayage des rayons gamma se compose d'un tube Geiger portable, d'une alimentation haute tension et d'un compteur de taux de comptage d'impulsions. La fréquence des impulsions est alors considérée comme une indication de l'intensité de l'exposition aux rayons gamma. Les échelles des compteurs de taux de comptage sont souvent calibrées en termes d'unités de taux d'exposition, mais dans certaines circonstances, ces lectures peuvent être erronées d'un facteur de 2 ou 3 ou plus.

La difficulté provient du fait que le taux de comptage d'un tube Geiger, contrairement au courant mesuré d'une chambre d'ionisation, n'a aucune relation fondamentale avec le taux d'exposition aux

rayons gamma. Pour une énergie de rayons gamma donnée, le taux de comptage évolue évidemment de manière linéaire avec l'intensité, mais les applications sont susceptibles d'impliquer des rayons gamma de nombreuses énergies différentes et variables. L'échelle du taux d'exposition peut être calibrée avec précision à n'importe quelle énergie fixe de rayons gamma, mais si l'appareil de mesure est appliqué à des mesures impliquant d'autres énergies de rayons gamma, la variation de l'efficacité du compteur avec l'énergie doit être prise en compte.

Le tube peut être recouvert d'une couche fine externe de métal tel que du plomb et / ou de l'étain, assurant ce que l'on appelle souvent une compensation énergétique du tube. Un tube G-M nu aura tendance à avoir une efficacité trop élevée à basse énergie par rapport à des énergies plus élevées pour une bonne correspondance avec la courbe d'exposition en fonction de l'énergie, et l'absorption préférentielle du filtre métallique à des énergies plus basses fournit une sorte de compensation empirique. Le rapport des deux courbes à n'importe quelle énergie donnée est une mesure du facteur de correction qui doit être appliqué aux relevés des compteurs.

Annexe II

Détecteurs à scintillation (scintillateurs)

1. Principes des détecteurs à scintillation

La détection des rayonnements ionisants par la lumière de scintillation produite dans certains matériaux est l'une des plus anciennes techniques connues. Le processus de scintillation reste l'une des méthodes les plus utiles disponibles pour la détection et la spectroscopie d'un large éventail de rayonnements. Dans cette annexe, nous discutons des différents types de scintillateurs disponibles et des considérations importantes dans la collecte efficace de la lumière de scintillation. Par la suite, les capteurs de lumière modernes, les tubes photomultiplicateurs et les photodiodes nécessaires pour convertir la lumière en une impulsion électrique, ainsi que l'application des détecteurs à scintillation en spectroscopie de rayonnement seront traités en détail.

Le matériau de scintillation idéal doit posséder les propriétés suivantes:

1. Il doit convertir l'énergie cinétique des particules chargées en lumière détectable avec une efficacité de scintillation élevée.
2. Cette conversion doit être linéaire - le rendement lumineux doit être proportionnel à l'énergie déposée sur une plage aussi large que possible.
3. Le support doit être transparent à la longueur d'onde de sa propre émission pour une bonne collecte de la lumière.
4. Le temps de décroissance de la luminescence induite doit être court afin que des impulsions rapides puissent être générées.
5. Le matériau doit être de bonne qualité optique et peut être fabriqué dans des tailles suffisamment grandes pour présenter un intérêt en tant que détecteur pratique.
6. Son indice de réfraction doit être proche de celui du verre (environ 1.5) pour permettre un couplage efficace de la lumière de scintillation à un tube photomultiplicateur ou à un autre capteur de lumière.

Aucun matériau ne répond simultanément à tous ces critères, et le choix d'un scintillateur particulier est toujours un compromis entre ces facteurs et d'autres. Les scintillateurs les plus largement appliqués comprennent les cristaux d'halogénures alcalins inorganiques, dont l'iodure de sodium est le préféré, les liquides et les plastiques à base organique. Les scintillateurs inorganiques ont tendance à avoir le meilleur rendement lumineux et la meilleure linéarité, mais à quelques exceptions près, leur temps de réponse est relativement lent. Les scintillateurs organiques sont généralement plus rapides mais produisent moins de lumière. L'application envisagée a également une influence majeure sur le choix du scintillateur. La valeur Z élevée des constituants et la densité

élevée des cristaux inorganiques favorisent leur choix pour la spectroscopie gamma, tandis que les composés organiques sont souvent préférés pour la spectroscopie bêta et la détection de neutrons rapides (en raison de leur teneur en hydrogène).

Le processus de fluorescence est l'émission rapide de rayonnement visible d'une substance suite à son excitation par différents moyens. Il est classique de distinguer plusieurs autres processus pouvant également conduire à l'émission de lumière visible. La phosphorescence correspond à l'émission d'une lumière de plus grande longueur d'onde que la fluorescence, et avec un temps caractéristique généralement beaucoup plus lent. La fluorescence retardée donne le même spectre d'émission que la fluorescence rapide mais se caractérise à nouveau par un temps d'émission beaucoup plus long après l'excitation. Pour être un bon scintillateur, un matériau doit convertir une fraction aussi grande que possible de l'énergie de rayonnement incident pour déclencher la fluorescence, tout en minimisant les effets généralement indésirables pour la phosphorescence et la fluorescence retardée. Cette annexe met principalement l'accent sur le fonctionnement en mode impulsionnel des scintillateurs. La lumière pouvant contribuer à une impulsion de sortie est généralement limitée à la fluorescence rapide car les temps de mise en forme des impulsions du circuit de mesure sont beaucoup plus faibles que les temps de phosphorescence typiques et de décroissance de la fluorescence retardée. Cette lumière à longue durée de vie est ensuite répartie plus ou moins aléatoirement dans le temps entre les impulsions du signal et arrive au capteur de lumière sous forme de photons individuels qui souvent ne peuvent pas être distingués du bruit aléatoire. En revanche, les scintillateurs qui fonctionnent en mode courant sous un éclairage constant produiront un signal de courant à l'état stable qui est proportionnel au rendement lumineux total, et toutes les composantes de désintégration devront contribuer d'un taux à leur intensité absolue. Pour cette raison, le rendement lumineux mesuré à partir d'un scintillateur fonctionnant en mode impulsionnel peut apparaître inférieur à celui déduit d'un fonctionnant en mode courant de régime permanent enregistré à partir du même scintillateur. Les détecteurs à scintillation fonctionnant en mode courant dans des conditions dans lesquelles l'intensité du rayonnement change rapidement souffriront d'effets de mémoire ou de "rémanence" si les composants de désintégration à longue durée de vie sont importants.

Dans cette annexe, nous limitons les discussions aux processus nécessaires pour comprendre les différences de comportement des différents types de scintillateurs, ainsi que certaines de leurs propriétés importantes en tant que détecteurs de rayonnement pratiques.

2. Scintillateurs organiques

2.1 Mécanisme de scintillation dans les composés organiques

Le processus de fluorescence dans les matériaux organiques découle des transitions dans la structure du niveau d'énergie d'une seule molécule et peut donc être observé à partir d'une espèce moléculaire donnée indépendamment de son état physique. Par exemple, on observe que l'anthracène, sous forme de matériau polycristallin solide, sous forme de vapeur ou dans le cadre d'une solution à plusieurs composants émet de la fluorescence. Ce comportement est en contraste marqué avec les scintillateurs inorganiques cristallins tels que l'iodure de sodium, qui nécessitent un réseau cristallin régulier comme base pour le processus de scintillation.

Une grande catégorie de scintillateurs organiques pratiques est basée sur des molécules organiques avec certaines propriétés de symétrie qui donnent naissance à ce que l'on appelle une structure de π électrons. Les niveaux d'énergie π électronique d'une telle molécule sont illustrés à la figure (01). L'énergie peut être absorbée en excitant la configuration électronique dans un état quelconque de l'un des états excités. Une série d'états singulet (spin 0) sont désignés par : $S_0, S_1, S_2, \dots$ dans la figure. Un ensemble similaire de niveaux électroniques triplet (spin 1) est également représenté par $T_1, T_2, T_3, \dots$. Pour les molécules d'intérêt en tant que scintillateurs organiques, l'espacement d'énergie entre S_0 et S_1 est de 3 ou 4 eV, alors que l'espacement entre les états « higher-lying » sont généralement un peu plus petits. Chacune de ces configurations électroniques est en outre subdivisée en une série de niveaux beaucoup plus espacés qui correspondent à divers états vibrationnels de la molécule. L'espacement typique de ces niveaux est de l'ordre de 0,15 eV. Un deuxième indice est souvent ajouté pour distinguer ces états vibrationnels, et le symbole S_{00} représente l'état vibrationnel le plus bas de l'état électronique fondamental.

Comme l'espacement entre les états vibrationnels est important par rapport aux énergies thermiques moyennes (0,025 eV), presque toutes les molécules à température ambiante sont à l'état S_{00} . Sur la figure (01), l'absorption d'énergie par la molécule est représentée par les flèches pointant vers le haut. Dans le cas d'un scintillateur, ces processus représentent l'absorption de l'énergie cinétique d'une particule chargée passant à proximité. Les états électroniques singulets supérieurs qui sont excités sont rapidement (de l'ordre des picosecondes) désexcités à l'état électronique S_1 par conversion interne sans rayonnement. De plus, tout état avec un excès d'énergie vibratoire (tel que S_{11} ou S_{12}) n'est pas en équilibre thermique avec ses voisins et perd à nouveau rapidement cette énergie vibrationnelle. Par conséquent, l'effet net du processus d'excitation dans un cristal organique simple est de produire, après une période de temps négligeable, une population de molécules excitées à l'état S_{10} .

La lumière de scintillation principale (ou fluorescence prompte) est émise dans les transitions entre cet état S_{10} et l'un des états vibrationnels de l'état électronique fondamental. Ces transitions sont indiquées par les flèches vers le bas sur la figure (01). Si τ représente le temps de décroissance de la fluorescence pour le niveau S_{10} , alors l'intensité de fluorescence rapide à un instant t suivant l'excitation doit être :

$$I = I_0 e^{-t/\tau} \quad (01)$$

Dans la plupart des scintillateurs organiques, τ est de quelques nanosecondes, et la composante de scintillation prompte est donc relativement rapide. La durée de vie du premier état triplet T_1 est de manière caractéristique beaucoup plus longue que celle de l'état singulet S_1 . Grâce à une transition appelée croisement inter-système, certains états singulets excités peuvent être convertis en états triplets. La durée de vie de T_1 peut atteindre 10^{-3} s et le rayonnement émis lors d'une désexcitation de T_1 à S_0 est donc une émission lumineuse retardée caractérisée comme une phosphorescence. Puisque T_1 se situe en dessous de S_1 , la longueur d'onde de ce spectre de phosphorescence sera plus longue que celle du spectre de fluorescence. Pendant qu'elles sont dans l'état T_1 , certaines molécules peuvent être excitées thermiquement pour revenir à l'état S_1 et ensuite se désintégrer par fluorescence normale. Ce processus représente l'origine de la fluorescence retardée parfois observée pour les matières organiques.

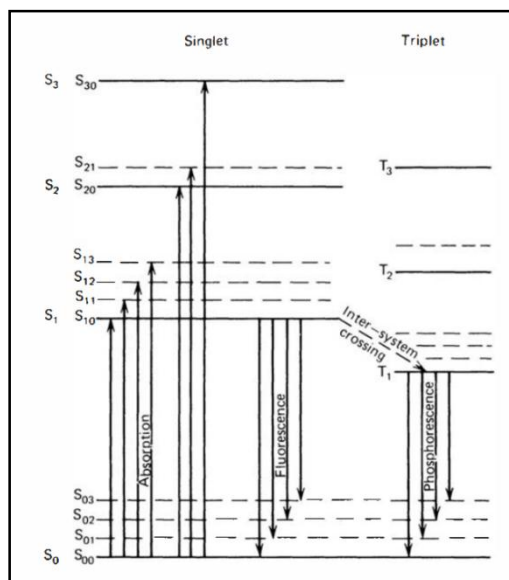


Fig. 01 Niveaux d'énergie d'une molécule organique avec une structure π électrons

La figure (01) peut être également utilisée pour expliquer pourquoi les scintillateurs organiques peuvent être transparents à leur propre émission de fluorescence. La longueur des flèches vers le haut

correspond aux énergies des photons qui seront fortement absorbées dans le matériau. Parce que toutes les transitions de fluorescence représentées par les flèches vers le bas (à l'exception de $S_{10}-S_{00}$) ont une énergie inférieure au minimum requis pour l'excitation, il y a très peu de chevauchement entre les spectres d'absorption et d'émission optiques (souvent appelé décalage de Stokes), et par conséquent il en résulte peu d'auto-absorption de la fluorescence. Un exemple de ces spectres pour un scintillateur organique typique est donné à la figure (02).

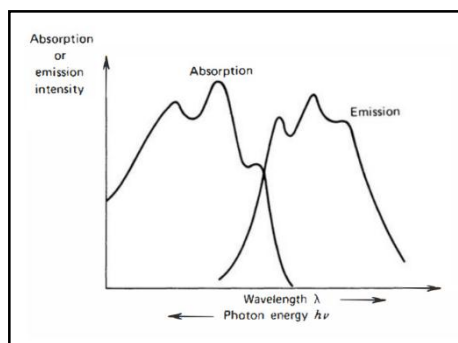


Fig. 02 Spectres d'absorption et d'émission optiques d'un scintillateur organique typique avec la structure de niveau illustrée à la figure (01)

L'efficacité de scintillation de tout scintillateur est définie comme le taux de toute l'énergie des particules incidentes qui est convertie en lumière visible. Il est préféré toujours que ce rendement soit le plus grand possible, mais il existe malheureusement des modes alternatifs de désexcitation possibles pour les molécules excitées qui n'impliquent pas l'émission de lumière et dans lesquels l'excitation se dégrade principalement en chaleur. Tous ces processus de désexcitation sans rayonnement sont regroupés sous le terme d'extinction. Dans la fabrication et l'utilisation de scintillateurs organiques, il est souvent important d'éliminer les impuretés (telles que l'oxygène dissous dans les scintillateurs liquides), qui dégradent le rendement lumineux en fournissant des mécanismes d'extinction alternatifs pour l'énergie d'excitation.

Dans presque tous les matériaux organiques, l'énergie d'excitation subit un transfert substantiel de molécule à molécule avant que la désexcitation ne se produise. Ce processus de transfert d'énergie est particulièrement important pour la grande catégorie de scintillateurs organiques qui implique plus d'une espèce de molécules. Si une petite concentration d'un scintillateur efficace est ajoutée à un solvant, l'énergie qui est absorbée, principalement par le solvant, peut éventuellement trouver son chemin vers l'un des molécules de scintillation efficaces et provoquent une émission de lumière. Ces scintillateurs organiques "binaires" sont largement utilisés à la fois comme solutions liquides et plastiques incorporant une variété de solvants et de scintillants organiques dissous.

Un troisième composant est parfois ajouté à ces mélanges pour décaler la longueur d'onde. Sa fonction est d'absorber la lumière produite par le scintillant primaire et de la rediffuser à une plus grande longueur d'onde. Ce décalage dans le spectre d'émission peut être utile pour une correspondance plus étroite avec la sensibilité spectrale d'un tube photomultiplicateur ou pour minimiser l'auto-absorption dans les grands scintillateurs liquides ou plastiques. Un passage en revue a été établi des mécanismes de transfert d'énergie en binaire et mélanges organiques tertiaires ainsi que leur influence sur l'efficacité de la scintillation et les caractéristiques de synchronisation des impulsions.

2.2 Types de scintillateurs organiques

2.2.1 Cristaux organiques purs

Seuls deux matériaux ont acquis une large popularité en tant que scintillateurs cristallins organiques purs. L'Anthracène est l'un des plus anciens matériaux organiques utilisés en scintillation et a la particularité d'avoir l'efficacité de scintillation la plus élevée (ou la plus grande puissance lumineuse par unité d'énergie) de tous les scintillateurs organiques. Le Stilbène a une efficacité de scintillation inférieure mais il est préféré dans les situations où la discrimination de forme d'impulsion doit être utilisée pour distinguer les scintillations induites par des particules chargées et des électrons. Les deux matériaux sont relativement fragiles et difficiles à obtenir en grandes tailles. En outre, l'efficacité de scintillation est connue pour dépendre de l'orientation d'une particule ionisante par rapport à l'axe du cristal. Cette variation directionnelle, qui peut atteindre 20 à 30 %, gâche la résolution énergétique pouvant être obtenue avec ces cristaux si le rayonnement incident produit des traces dans diverses directions à l'intérieur du cristal. Ces monocristaux sont limités à pas plus d'environ 10 cm de diamètre et d'épaisseur maximum. Il y a eu quelques progrès dans le développement du Stilbène polycristallin par pressage à chaud ou à froid de petits grains du matériau pour produire des détecteurs de plus grand diamètre (jusqu'à 25 cm) mais d'épaisseur plus limitée.

2.2.2 Solutions liquides organiques

Une catégorie de scintillateurs utiles est produite en dissolvant un scintillateur organique dans un solvant approprié. Les scintillateurs liquides peuvent être constitués simplement de ces deux composants, ou un troisième composant est parfois ajouté pour décaler la longueur d'onde pour adapter le spectre d'émission à mieux correspondre à la réponse spectrale des tubes photomultiplicateurs courants. Les scintillateurs liquides sont souvent renfermés dans des récipients en verre scellés et sont manipulés de la même manière que les scintillateurs solides. Dans certaines applications, des détecteurs de grand volume avec des dimensions de plusieurs mètres peuvent être nécessaires. Dans ces cas, le scintillateur liquide est souvent le seul choix pratique du point de vue du

coût. Dans de nombreux liquides, la présence d'oxygène dissous peut servir d'agent d'extinction puissant et peut conduire à une efficacité de fluorescence considérablement réduite. Il est alors nécessaire que la solution soit enfermée dans un volume clos dont l'essentiel de l'oxygène a été purgé. En raison de leur absence de structure solide qui pourrait être endommagée par une exposition à un rayonnement intense, les scintillateurs liquides devraient être plus résistants aux effets des dommages causés par les rayonnements que les scintillateurs cristallins ou plastiques. Cette attente est confirmée par des mesures, et une résistance raisonnable au changement jusqu'à des expositions aussi élevées que 10^5 Gy a été démontrée pour certains liquides.

Les scintillateurs liquides sont également largement utilisés pour compter les matières radioactives qui peuvent être dissoutes dans la solution de scintillateur. Dans ce cas, tous les rayonnements émis par la source traversent immédiatement une partie du scintillateur et le rendement de comptage peut être proche de 100 %. La technique est largement utilisée pour compter l'activité bêta de bas niveau, comme celle du carbone 14 ou du tritium.

2.2.3 Scintillateurs en plastique

Si un scintillateur organique est dissous dans un solvant qui peut ensuite être polymérisé ultérieurement, l'équivalent d'une solution solide peut être produit. Un exemple courant est un solvant constitué de monomère de Styrène dans lequel un scintillateur organique approprié est dissous. Le Styrène est ensuite polymérisé pour former un plastique solide. D'autres matrices plastiques peuvent être constituées de polyvinyltoluène ou de polyméthacrylate de méthyle. En raison de la facilité avec laquelle ils peuvent être façonnés et fabriqués, les plastiques sont devenus une forme extrêmement utile de scintillateur organique.

Les scintillateurs en plastique sont disponibles avec une bonne sélection des tailles standards de barreaux, de cylindres et de feuilles plates. Comme le matériau est relativement peu coûteux, les plastiques sont souvent le seul choix pratique si des scintillateurs solides de grand volume sont nécessaires. Dans ces cas, l'auto-absorption de la lumière du scintillateur peut ne plus être négligeable, et une certaine attention doit être accordée aux propriétés d'atténuation du matériau. La distance sur laquelle l'intensité lumineuse sera atténuée d'un facteur 2 peut atteindre plusieurs mètres, bien que des longueurs d'atténuation beaucoup plus faibles soient observées pour certains plastiques.

Il existe également une large sélection de scintillateurs en plastique disponibles sous forme de fibres de petit diamètre. Utilisés sous forme de fibres simples ou regroupés sous forme de faisceaux ou de rubans, ces scintillateurs se prêtent à des applications dans lesquelles la position des interactions de particules doit être détectée avec une bonne résolution spatiale.

En raison de l'application généralisée des scintillateurs plastiques dans les mesures de la physique des particules, où ils peuvent être exposés à des niveaux de rayonnement élevés et soutenus, une attention

considérable a été accordée à la dégradation de la production de scintillation des plastiques due aux dommages causés par les rayonnements. Ce processus est complexe et de nombreuses variables telles que le débit de dose, la présence ou l'absence d'oxygène et la nature du rayonnement jouent un rôle important. Il existe également une tendance à observer une certaine récupération ou un recuit des dommages sur des périodes de temps qui peuvent être des heures ou des jours après une exposition. Dans les scintillateurs plastiques typiques, une dégradation significative du rendement lumineux est observée pour des expositions cumulées aux rayons gamma dans la plage de 10^3 ou 10^4 Gy, alors que d'autres formulations résistantes aux rayonnements montrent peu de diminution du rendement lumineux avec des doses aussi élevées que 10^5 Gy. Les changements dans la lumière mesurée peuvent consister soit en une diminution du rendement lumineux causée par des dommages au composant fluorescent, soit en une diminution de la transmission de la lumière causée par la création de centres d'absorption optique.

2.2.4 Scintillateurs à couches minces

Les couches très minces de scintillateur plastique jouent un rôle unique dans le domaine des détecteurs de rayonnement. Etant donné que des films ultrafins d'une épaisseur aussi faible que $20 \mu\text{g}/\text{cm}^2$ peuvent être fabriqués, il est facile de fournir un détecteur qui est mince par rapport à la gamme de particules même faiblement pénétrantes telles que les ions lourds. Ces films servent ainsi de détecteurs de transmission, qui ne répondent qu'au taux d'énergie perdue par la particule lorsqu'elle traverse le détecteur. L'épaisseur peut être jusqu'à un ou deux ordres de grandeur plus petite que le minimum possible avec d'autres configurations de détecteur, telles que des barrières de surface en silicium totalement appauvries. Les films sont disponibles dans le commerce avec des épaisseurs jusqu'à environ $10 \mu\text{m}$. Des films encore plus minces peuvent être produits par l'utilisateur grâce à des techniques telles que l'évaporation à partir d'une solution de scintillateur plastique ou par le procédé de revêtement par centrifugation. Le film peut être déposé directement sur la face d'un tube photomultiplicateur, ou la lumière peut être collectée indirectement par un conduit de lumière transparent en contact avec les bords du film. Alternativement, le film peut être placé dans une cavité réfléchissante. La réponse de ces films ne découle pas directement de la perte d'énergie attendue des ions dans le détecteur et est une fonction plus complexe de la vitesse des ions et du numéro atomique. Le rendement lumineux par unité d'énergie perdue augmente avec la diminution du numéro atomique de l'ion, de sorte que les films minces peuvent être des détecteurs de transmission utiles pour les protons ou les particules alpha même lorsque l'énergie déposée est relativement faible. Comme d'autres scintillateurs organiques, les détecteurs à couches minces présentent des temps de décroissance de la scintillation de quelques ns seulement, et ils se sont révélés très utiles dans les mesures temporelles rapides.

2.2.5 Scintillateurs organiques chargés

Les scintillateurs organiques en tant que catégorie sont généralement utiles pour la détection directe de particules « bêta » (électrons rapides) ou de particules alpha (ions positifs). Ils sont également facilement adaptables à la détection de neutrons rapides par le processus de recul des protons. Cependant, en raison de la faible valeur Z de leurs constituants (hydrogène, carbone et oxygène), il n'y a pratiquement pas de section efficace photoélectrique pour les rayons gamma d'énergies typiques. En conséquence, les scintillateurs organiques typiques ne présentent aucun photo-pic et ne donneront lieu qu'à un continuum Compton dans leur spectre de hauteur d'impulsion de rayons gamma.

Pour fournir un certain degré de conversion photoélectrique des rayons gamma, des tentatives ont été faites pour ajouter des éléments à Z élevé aux scintillateurs organiques. La forme la plus courante est l'ajout de plomb ou d'étain à des scintillateurs plastiques courants jusqu'à une concentration de 10 % en poids. Il a été également démontré que l'étain peut être ajouté à des solutions de scintillateur organique liquide à des concentrations allant jusqu'à 54 % en poids tout en conservant un rendement lumineux de scintillation quelque peu diminué. Aux basses énergies des rayons gamma, l'efficacité du photo-pic de ces matériaux peut être rendue relativement élevée, et ils présentent les avantages supplémentaires d'une réponse rapide et d'un faible coût par rapport aux scintillateurs à rayons gamma plus conventionnels. Malheureusement, l'ajout de ces éléments à Z élevé conduit inévitablement à une diminution du rendement lumineux, et la résolution énergétique qui peut être obtenue est donc considérablement inférieure à celle des scintillateurs inorganiques.

D'autres exemples de chargement de scintillateurs organiques surviennent en relation avec la détection de neutrons. Les scintillateurs liquides ou plastiques peuvent être ensemencés avec l'un des éléments à forte section efficace pour les neutrons tels que le bore, le lithium ou le gadolinium. Les particules chargées secondaires et/ou les rayons gamma produits par les réactions induites par les neutrons peuvent ensuite être détectés directement dans le scintillateur pour fournir un signal de sortie.

2.3 Réponse des scintillateurs organiques

2.3.1 Lumière de sortie

Une petite fraction de l'énergie cinétique perdue par une particule chargée dans un scintillateur est convertie en énergie fluorescente. Le reste est dissipé de manière non radiative, principalement sous la forme de vibrations de réseau ou de chaleur. La fraction de l'énergie des particules qui est convertie (l'efficacité de scintillation) dépend à la fois du type de particule et de son énergie. Dans certains cas, l'efficacité de scintillation peut être indépendante de l'énergie, conduisant à une dépendance linéaire du rendement lumineux sur l'énergie initiale.

Pour les scintillateurs organiques tels que l'Anthracène, le Stilbène et de nombreux scintillateurs liquides et plastiques disponibles, la réponse aux électrons est linéaire pour les énergies de particules supérieures à environ 125 ke V. La réponse aux particules chargées lourdes telles que les protons ou les particules alpha est toujours moindre pour des énergies équivalentes et est non linéaire à des énergies initiales beaucoup plus élevées.

La réponse des scintillateurs organiques aux particules chargées peut être mieux décrite par une relation entre dL / dx , l'énergie fluorescente émise par unité de longueur de trajet, et dE / dx , la perte d'énergie spécifique pour la particule chargée. Une autre hypothèse est qu'en l'absence d'extinction, le rendement lumineux est proportionnel à la perte d'énergie :

$$\boxed{\frac{dL}{dx} = S \frac{dE}{dx}} \quad (02)$$

Où S est le rendement normal de scintillation. Pour tenir compte de la probabilité d'extinction.

Lorsqu'il est excité par des électrons rapides (soit directement, soit par irradiation gamma), dE / dx est petit pour des valeurs suffisamment grandes de E et la formule (02) prédit alors :

$$\boxed{\left. \frac{dL}{dx} \right|_e = S \frac{dE}{dx}} \quad (03)$$

Où le flux lumineux supplémentaire par unité de perte d'énergie est une constante :

$$\boxed{\left. \frac{dL}{dE} \right|_e = S} \quad (04)$$

C'est le régime dans lequel le flux lumineux :

$$\boxed{L \equiv \int_0^E \frac{dL}{dE} dE = SE} \quad (05)$$

est linéairement lié à l'énergie particulaire initiale E.

D'autre part, pour une particule alpha, dE / dx est très grand de sorte que la saturation se produit le long de la piste et la formule (02) devient :

$$\left. \frac{dL}{dx} \right|_{\alpha} = \frac{S}{kB}$$

(06)

Le rapport alpha / bêta est un paramètre largement utilisé pour décrire la différence de rendement lumineux d'un scintillateur organique pour les électrons et les particules chargées de même énergie. Le rendement lumineux des électrons est toujours supérieur à celui d'une particule chargée de même énergie cinétique, et donc le rapport alpha / bêta est toujours inférieur à 1. Ce rapport dépendra de l'énergie à laquelle la comparaison est faite, et aucune valeur fixe n'est applicable sur toute la plage d'énergie.

Dans certains composés organiques, le chevauchement partiel des spectres d'émission et d'absorption conduit à une dépendance de la taille de l'efficacité apparente pour la scintillation. De plus, la réponse anisotrope de l'anthracène et du Stilbène complique davantage les mesures d'efficacité. Une exposition prolongée aux rayonnements ionisants conduit à une détérioration générale des propriétés de la plupart des scintillateurs organiques. Il a été également démontré que les scintillateurs plastiques exposés à la lumière et à l'oxygène subissent une détérioration à long terme due à la dégradation des polymères. En plus des effets internes, la surface des plastiques peut souvent se fissurer lors d'une exposition à des environnements extrêmes. La fissuration de la surface entraîne une baisse substantielle du flux lumineux observé des grands scintillateurs en raison de la diminution de l'efficacité réflexion interne de la lumière.

2.3.2 Temps de réponse

Si l'on peut supposer que les états luminescents dans une molécule organique se forment instantanément et que seule une fluorescence rapide est observée, alors le profil temporel de l'impulsion lumineuse devrait être un front montant très rapide suivi d'une simple décroissance exponentielle. Bien que cette représentation simple soit souvent adéquate pour de nombreuses descriptions du comportement du scintillateur, un modèle plus détaillé de la dépendance temporelle du rendement de scintillation doit prendre en compte deux autres effets : le temps fini nécessaire pour peupler les états luminescents et les composantes plus lentes de la scintillation correspondant à la fluorescence et à la phosphorescence retardées.

Des temps d'environ une demi-nanoseconde sont nécessaires pour peupler les niveaux à partir desquels apparaît la lumière de fluorescence rapide. Pour les scintillateurs très rapides, le temps de décroissance à partir de ces niveaux n'est que trois ou quatre fois supérieur, et une description complète de la forme d'impulsion attendue doit également tenir compte du temps de montée fini. Une approche suppose que la population des niveaux optiques est également exponentielle et que la forme globale de l'impulsion lumineuse est donnée par :

$$I = I_0 \left(e^{-t/\tau} - e^{-t/\tau_1} \right) \quad (07)$$

Où τ_1 est la constante de temps décrivant la population des niveaux optiques et τ est la constante de temps décrivant leur décroissance. Les valeurs de ces paramètres pour plusieurs scintillateurs plastiques sont données dans le tableau (01). D'autres observations ont conclu que le pas de population est mieux représenté par une fonction gaussienne $f(t)$ caractérisée par un écart-type σ_{ET} . Le profil global de lumière en fonction du temps est alors décrit par :

$$\frac{I}{I_0} = f(t) e^{-t/\tau} \quad (08)$$

Les meilleures valeurs d'ajustement pour σ_{ET} sont également présentées dans le tableau (02). Expérimentalement, la montée et la descente de la sortie lumineuse peuvent être caractérisées par la pleine largeur à mi-hauteur (FWHM) du profil de lumière en fonction du temps résultant, qui peut être mesuré à l'aide de procédures de synchronisation très rapides. Il est devenu courant de spécifier les performances des scintillateurs organiques ultrarapides par leur temps FWHM plutôt que par le temps de décroissance seul.

Tab. 01 Quelques propriétés de synchronisation des scintillateurs plastiques rapides

Some Timing Properties of Fast Plastic Scintillators					
	Parameters for Eq. (10)		Parameters for Eq. (11)		Measured FWHM
	τ_1 (rise)	τ (decay)	σ_{ET}	τ	
NE 111	0.2 ns	1.7 ns	0.2 ns	1.7 ns	1.54 ns
Naton 136	0.4 ns	1.6 ns	0.5 ns	1.87 ns	2.3 ns
NE 102A	0.6 ns	2.4 ns	0.7 ns	2.4 ns	3.3 ns

Pour les scintillateurs très rapides, des effets autres que le mécanisme de scintillation peuvent parfois affecter la réponse temporelle observée. Parmi ceux-ci figure le temps de vol fini des photons du point de scintillation au tube photomultiplicateur. En particulier dans les grands scintillateurs, la fluctuation du temps de transit due aux multiples réflexions de la lumière sur les surfaces du scintillateur peut équivaloir à une propagation importante du temps d'arrivée de la lumière à la photocathode du tube photomultiplicateur.

En outre, il existe des preuves que l'auto-absorption et la réémission de la fluorescence jouent un rôle important en provoquant une détérioration apparente de la résolution temporelle lorsque les dimensions d'un scintillateur sont agrandies.

Bien qu'un temps de décroissance rapide soit un avantage dans presque toutes les applications des scintillateurs, il existe au moins une application dans laquelle une décroissance lente est nécessaire.

2.3.3 Discrimination en forme d'impulsions

Pour la grande majorité des scintillateurs organiques, la fluorescence instantanée représente la majeure partie de la lumière de scintillation observée. Une composante à durée de vie plus longue est également observée dans de nombreux cas, correspondant toutefois à une fluorescence retardée. La courbe du rendement composite peut être souvent représentée de manière adéquate par la somme de deux décroissances exponentielles, appelées composantes rapide et lente de la scintillation. Par rapport au temps de décroissance rapide de quelques nanosecondes, la composante lente aura typiquement un temps de décroissance caractéristique de plusieurs centaines nanosecondes. Étant donné que la majorité du rendement lumineux se produit dans la composante rapide, la queue à longue durée de vie n'aurait pas une grande conséquence, à l'exception d'une propriété très utile : la fraction de lumière qui apparaît dans la composante lente dépend souvent de la nature de la particule excitante. On peut donc se servir de cette dépendance pour différencier des particules de natures différentes qui déposent la même énergie dans le détecteur. Ce processus est souvent appelé discrimination de forme d'impulsion et est largement appliqué pour éliminer les événements induits par les rayons gamma lorsque des scintillateurs organiques sont utilisés comme détecteurs de neutrons.

Il existe des preuves solides que la composante de scintillation lente provient de l'excitation d'états triplets à longue durée de vie (appelés T_1 sur la figure 01) le long de la trajectoire de la particule ionisante.

Les interactions bimoléculaires entre deux de ces molécules excitées peuvent produire des molécules, l'une dans l'état singulet le plus bas (S_1) et l'autre dans l'état fondamental. La molécule à l'état singulet peut alors se désexciter de manière normale, entraînant une fluorescence retardée. La variation du rendement du composant lent peut alors s'expliquer en partie par les différences attendues dans la densité des états triplets le long de la trajectoire de la particule, car le rendement de la réaction bimoléculaire devrait dépendre du carré de la concentration du triplet. Par conséquent, la fraction de composante lente devrait dépendre principalement du taux de perte d'énergie dE / dx de la particule excitatrice et devrait être la plus grande pour les particules avec un grand dE / dx . Ces prédictions sont généralement confirmées par des mesures de la forme des impulsions de scintillation à partir d'une grande variété de matières organiques.

Certains scintillateurs organiques, y compris les cristaux de Stilbène et un certain nombre de scintillateurs liquides commerciaux, sont particulièrement appréciés pour la discrimination de la forme des impulsions en raison des grandes différences dans la composante lente relative induite par

différents rayonnements. La figure (03) montre les différences observées dans le Stilbène pour les particules alpha, les neutrons rapides (protons de recul) et les rayons gamma (électrons rapides). Dans de tels scintillateurs, il est non seulement possible de différencier les rayonnements avec de grandes différences dE/dx (telles que les neutrons et les rayons gamma) mais aussi pour séparer les événements résultant de diverses espèces de particules chargées lourdes. Des revues des propriétés de discrimination de la forme des impulsions de différents types de scintillateurs organiques et des exemples d'applications sont données dans plusieurs publications. Les techniques électroniques pour effectuer cette discrimination de forme d'impulsion sont de plus en plus basées sur la numérisation de la forme d'onde d'impulsion.

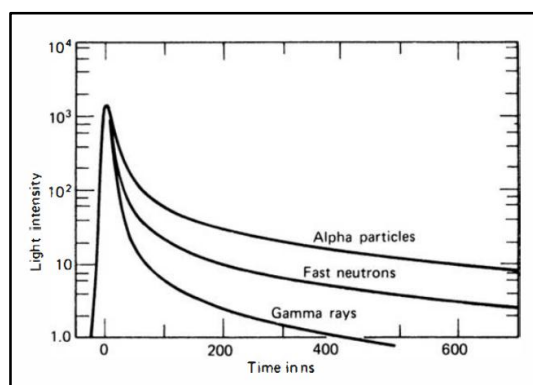


Fig. 03 Dépendance temporelle des impulsions de scintillation dans le Stilbène (Intensité égale au temps zéro) lorsqu'elles sont excitées par des radiations de différents types

3. Scintillateurs inorganiques

Après des décennies d'évolution relativement lente, le monde des scintillateurs inorganiques a connu une renaissance qui a commencé au milieu des années 1980 et se poursuit jusqu'à nos jours. Stimulés par la découverte surprenante de certains nouveaux matériaux qui offrent d'excellents rendements lumineux et/ou des temps de décroissance rapides, un grand nombre de matériaux de scintillation potentiels ont récemment fait l'objet de tests et d'évaluations dans des laboratoires du monde entier. Dans les sections qui suivent, nous nous concentrerons sur un nombre limité d'entre eux qui présentent actuellement les caractéristiques les plus attrayantes, ainsi que sur certains piliers disponibles depuis les années 1950. Étant donné que la recherche de nouveaux scintillateurs inorganiques reste assez active, la littérature actuelle doit être recherchée pour les derniers développements.

3.1 Mécanisme de scintillation dans les cristaux inorganiques avec des activateurs

Le mécanisme de scintillation dans les matériaux inorganiques dépend des états d'énergie déterminés par le réseau cristallin du matériau. Comme le montre la figure (04), les électrons ne disposent que de bandes d'énergie discrètes dans les matériaux classés comme isolants ou semi-conducteurs. La bande inférieure, appelée bande de valence, dans laquelle se trouvent les électrons qui sont essentiellement liés aux sites du réseau, tandis que dans la bande de conduction se trouvent les électrons qui ont suffisamment d'énergie pour être libres de migrer à travers le cristal. Il existe une bande intermédiaire d'énergies, appelée bande interdite, dans laquelle les électrons ne peuvent jamais se trouver dans le cristal pur. L'absorption d'énergie peut entraîner l'élévation d'un électron de sa position normale dans la bande de valence à travers l'espace dans la bande de conduction, laissant un trou dans la bande de valence normalement remplie. Dans le cristal pur, le retour de l'électron dans la bande de valence avec l'émission d'un photon est un processus inefficace.

De plus, les largeurs d'intervalle typiques sont telles que le photon résultant aurait une énergie trop élevée pour se situer dans la plage visible.

Pour augmenter la probabilité d'émission de photons visibles pendant le processus de désexcitation, de petites quantités d'une impureté sont couramment ajoutées aux scintillateurs inorganiques. Ces impuretés délibérément ajoutées, appelées activateurs, créent des sites spéciaux dans le réseau au niveau desquels la structure normale de la bande d'énergie est modifiée par rapport à celle du cristal pur. En conséquence, des états d'énergie seront créés dans l'espace interdit à travers lesquels l'électron peut se désexciter vers la bande de valence. Comme l'énergie est inférieure à celle de l'intervalle interdit complet, cette transition peut maintenant donner naissance à un photon visible et donc servir de base au processus de scintillation. Ces sites de désexcitation sont appelés centres de luminescence ou centres de recombinaison. Leur structure énergétique dans le réseau cristallin hôte détermine le spectre d'émission du scintillateur.

Une particule chargée traversant le milieu de détection formera un grand nombre de paires électron-trou créées par l'élévation des électrons de la bande de valence à la bande de conduction. Le trou positif dérivera rapidement vers l'emplacement d'un site activateur et l'ionisera, car l'énergie d'ionisation de l'impureté sera inférieure à celle d'un site de réseau typique. Pendant ce temps, l'électron est libre de migrer à travers le cristal et le fera jusqu'à ce qu'il rencontre un tel activateur ionisé. À ce stade, l'électron peut tomber dans le site de l'activateur, créant une configuration neutre qui peut avoir son propre ensemble d'états d'énergie excités. Ces états sont illustrés à la figure (04) sous forme de lignes horizontales à l'intérieur de l'intervalle interdit. Si l'état activateur qui se forme est une configuration excitée avec une transition autorisée vers l'état fondamental, sa désexcitation se produira très rapidement et avec une forte probabilité pour l'émission d'un photon correspondant. Si

l'activateur est bien choisi, cette transition peut se situer dans le domaine visible. Les durées de vie typiques pour de tels états excités sont de l'ordre de 30 à 500 ns. Étant donné que le temps de migration de l'électron est beaucoup plus court, toutes les configurations d'impuretés excitées se forment essentiellement en même temps et se désexciteront ensuite avec la caractéristique de demi-vie de l'état excité. C'est le temps de décroissance de ces états qui détermine donc les caractéristiques temporelles de la lumière de scintillation émise.

Certains scintillateurs inorganiques peuvent être caractérisés de manière adéquate par un seul temps de décroissance ou une simple exponentielle, bien qu'un comportement temporel plus complexe soit souvent observé.

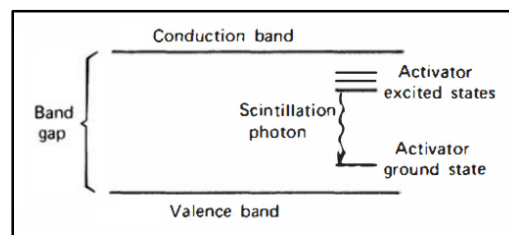


Fig. 04 Structure de bandes d'énergie d'un scintillateur cristallin activé

Il existe des procédés qui concurrencent celui qui vient d'être décrit. Par exemple, l'électron en arrivant sur le site de l'impureté peut créer une configuration excitée dont le passage à l'état fondamental est interdit. De tels états nécessitent alors un incrément supplémentaire d'énergie pour les élever à un état supérieur à partir duquel une désexcitation vers l'état fondamental est possible. Une source de cette énergie est l'excitation thermique, et la composante lente de la lumière qui en résulte est appelée phosphorescence. Il peut souvent être une source importante de lumière de fond ou de "rémanence" dans les scintillateurs.

Une troisième possibilité existe lorsqu'un électron est capturé dans un site activateur. Certaines transitions non radiatives sont possibles entre certains états excités formés par capture d'électrons et l'état fondamental, dans un tel cas aucun photon visible n'en résulte. De tels processus sont appelés « quenching » et représentent des mécanismes de perte dans la conversion de l'énergie des particules en lumière de scintillation.

Comme alternative à la migration indépendante de l'électron et du trou décrite ci-dessus, la paire peut à la place migrer ensemble dans une configuration faiblement associée connue sous le nom d'exciton. Dans ce cas, l'électron et le trou restent associés l'un à l'autre mais sont libres de dériver à travers le cristal jusqu'à atteindre le site d'un atome activateur. Des configurations d'activateur excité similaires peuvent à nouveau être formées et donner lieu à une lumière de scintillation dans leur désexcitation vers la configuration de masse.

Une mesure de l'efficacité du processus de scintillation découle d'un simple calcul d'énergie. Pour une large catégorie de matériaux, il faut en moyenne environ trois fois l'énergie de la bande interdite pour créer une paire électron-trou. Une conséquence importante de la luminescence à travers les sites activateurs est le fait que le cristal peut être transparent à la lumière de scintillation. Dans le cristal pur, il faudrait à peu près la même énergie pour exciter une paire électron-trou que celle libérée lorsque cette paire se recombine.

En conséquence, les spectres d'émission et d'absorption se chevauchent et il y aura une auto-absorption substantielle. Cependant, comme nous l'avons vu, l'émission d'un cristal activé se produit sur un site activateur où la transition d'énergie est inférieure à celle représentée par la création de la paire électron-trou. En conséquence, le spectre d'émission est décalé vers des longueurs d'onde plus grandes et ne sera pas influencé par la bande d'absorption optique de la masse du cristal.

De nombreux scintillateurs inorganiques (sinon la plupart) présentent plus d'un composant de désintégration, et des exemples dans lesquels un deuxième composant est important sont également indiqués. Lorsque des pourcentages sont indiqués, ils représentent les rendements relatifs des composants indiqués. La cinquième colonne est une estimation du nombre total de photons de scintillation produits sur tout le spectre d'émission à partir du dépôt de 1 Me V d'énergie par des électrons rapides. Ces valeurs sont généralement obtenues à partir de mesures dans lesquelles l'efficacité quantique du tube photomultiplicateur ou de la photodiode utilisée pour mesurer la lumière est connue en fonction de la longueur d'onde.

Pour les scintillateurs inorganiques courants, le rendement lumineux est plus proche de l'énergie de rayonnement déposée ce qui est généralement observé dans les scintillateurs organiques. Les processus d'extinction qui sont présents conduisent toujours à une certaine non-linéarité, mais souvent plus faible que dans les matières organiques.

On observe également une variation du rendement lumineux pour différents types de particules de même énergie. Comme dans les scintillateurs organiques, les particules chargées lourdes produisent moins de lumière par unité d'énergie. Le rapport alpha - bêta peut être beaucoup plus proche de l'unité que dans les composés organiques typiques, mais il est moins que 0,20 pour les matériaux à base d'oxyde.

3.2 Scintillateurs aux halogénures alcalins

- NaI (Tl) : Iodure de sodium avec une trace d'iodure de thallium ;
- CsI (Tl) et CsI (Na) : Iodure de césium avec du thallium ou du sodium ;
- LiI (Eu) : Iodure de lithium (activé à l'euporium).

3.3 Cristaux inorganiques lents (> 200 ns)

- GERMANATE DE BISMUTH (OU BGO) : $\text{B}_4\text{Ge}_3\text{O}_{12}$;
- CADMIUM TUNGSTATE (CdWO_4) et CALCIUM TUNGSTATE (CaWO_4) ;
- ZINC SULFIDE ZnS (Ag) : Sulfure de zinc activé à l'argent ;
- CALCIUM FLUORIDE CaF_2 (Eu) : Fluorure de calcium activé à l'euporium ;
- STRONTIUM IODIDE $\text{SrI}_2(\text{Eu})$: Iodure de strontium dopé à l'euporium $\text{SrI}_2(\text{Eu})$;

3.4 Cristaux inorganiques rapides non activés à faible rendement lumineux

- BARIUM FLUORIDE BaF_2 : Fluorure de baryum ;
- PURE CESIUM IODIDE CsI ;
- HALOGÉNURES DE CÉRIUM : CeF_3 , CeCl_3 , CeI_3 et CeBr_3 ;
- PLOMB TUNGSTATE PbWO_4 ;

3.5 Cristaux inorganiques "rapides" activés au cérium

- OXYORTHOSILICATE DE TERRES RARES ;
- PYROSILICATES LANTHANOÏDES ;
- PÉROVSKITES EN ALUMINIUM DE TERRES RARES ;
- GRENATS D'ALUMINIUM DE TERRES RARES ;
- HALOGÉNURES DE LANTHANE ;
- HALODURES DE LUTÉTIIUM ;
- ELPASOLITES ;

Scintillateurs en céramique transparents ;

Scintillateurs en verre ;

Scintillateurs au gaz nobles ;

Scintillateurs cryogéniques liquides et solides.

3.6 Effets des dommages causés par les rayonnements dans les scintillateurs inorganiques

Tous les matériaux de scintillation sont sujets aux effets des dommages causés par les rayonnements lorsqu'ils sont exposés à des flux élevés de rayonnement pendant de longues périodes. Les dommages sont très probablement mis en évidence par une réduction de la transparence du scintillateur causée par la création de centres de couleur qui absorbent la lumière scintillante. De plus, il peut également y avoir des interférences avec les processus qui donnent jusqu'à l'émission de la lumière de scintillation elle-même. Les expositions aux rayonnements peuvent également induire une

émission de lumière de longue durée sous forme de phosphorescence qui peut être gênante dans certains cas de mesure. On observe souvent que les effets des dommages dépendent du taux et varient fortement avec le type de rayonnement impliqué dans l'exposition. Les effets sont souvent au moins partiellement réversibles, le recuit ayant lieu même à température ambiante pendant des heures ou des jours suite à l'exposition.

4. Collecte de lumière et montage du scintillateur

4.1 Uniformité de la collecte de lumière

Dans tout détecteur à scintillation, on aimerait collecter la plus grande fraction possible de la lumière émise de façon isotrope à travers la trajectoire de la particule ionisante. Deux effets surviennent dans des cas pratiques qui conduisent à une collecte de lumière moins que parfaite : l'auto-absorption optique dans le scintillateur et les pertes aux surfaces du scintillateur. À l'exception des très grands scintillateurs (plusieurs centimètres de dimension) ou des matériaux de scintillation rarement utilisés (par exemple, ZnS), l'auto-absorption n'est généralement pas un mécanisme de perte significatif. Par conséquent, l'uniformité de la collecte de la lumière dépend principalement des conditions qui existent à l'interface entre le scintillateur et le conteneur dans lequel il est monté.

Les conditions de collecte de la lumière affectent la résolution en énergie d'un scintillateur en deux façons :

Premièrement, l'élargissement statistique de la fonction de réponse s'aggrave à mesure que le nombre de photons de scintillation qui contribuent à l'impulsion mesurée est réduit. La meilleure résolution ne peut donc être obtenue qu'en collectant la fraction maximale possible de tous les photons émis lors de l'événement de scintillation.

Deuxièmement, l'uniformité de la collection de lumière déterminera la variation de l'amplitude de l'impulsion du signal en tant que position de l'interaction de rayonnement varie dans tout le scintillateur. Une uniformité parfaite assurera que tous les événements déposant la même énergie, quel que soit l'endroit où ils se produisent dans le scintillateur, donnera lieu à la même amplitude moyenne des impulsions. Avec des scintillateurs ordinaires de quelques centimètres de dimensions, l'uniformité de la collecte de la lumière contribue rarement de manière significative à la résolution énergétique globale. Dans de grand scintillateur, en particulier ceux qui sont vus le long d'un bord mince, les variations de l'efficacité de la collecte de la lumière peut souvent dominer la résolution en énergie.

Puisque la lumière de scintillation est émise dans toutes les directions, seule une fraction limitée peut voyager directement sur la surface sur laquelle se trouve le tube photomultiplicateur ou un autre

capteur. Le reste, s'il doit être collecté, doit être réfléchi une ou plusieurs fois sur les surfaces du scintillateur.

4.2 Tubes de lumière

Il est souvent déconseillé voire impossible de coupler un tube photomultiplicateur directement à la face d'un scintillateur, car la taille ou la forme du scintillateur peut ne pas correspondre commodément à la surface circulaire de la photocathode des tubes PM disponibles. Une solution consiste à placer un tube de lumière entre le scintillateur mince et le tube PM qui diffusera la lumière de chaque événement de scintillation sur l'ensemble de la photocathode pour faire la moyenne de ces non-uniformités et améliorer la résolution de la hauteur d'impulsion.

Les tubes de lumière se basent sur le principe de la réflexion interne totale. Ce sont généralement des solides optiquement transparents avec un indice de réfraction relativement élevé pour minimiser l'angle critique pour la réflexion interne totale. Leurs surfaces sont très polies et sont souvent entourées d'une couche réfléchissante pour renvoyer une partie de la lumière qui s'échappe à des angles inférieurs à l'angle critique.

Pour maximiser la fraction de lumière qui est acheminée dans le tube de lumière, son indice de réfraction doit être aussi grand que possible. Cependant, la lumière est générée à l'intérieur du scintillateur, et non pas dans le tube de lumière, et c'est généralement l'indice de réfraction du scintillateur qui détermine la fraction de lumière collectée. Cela est particulièrement vrai lorsque le scintillateur est long dans la direction perpendiculaire à la surface de visualisation, et un photon de scintillation typique est réfléchi de manière multiple avant la collecte. Le scintillateur agit alors comme son propre conduit de lumière entre le point de scintillation et la surface de sortie. Comme le montre la figure (05), le couplage est réalisé à travers un certain nombre de bandes torsadées qui sont alignées au bord du scintillateur mais convergent vers un motif plus compact à l'extrémité du photomultiplicateur. L'unité peut facilement être fabriquée à partir de bandes de plastique plates, qui sont ensuite pliées après chauffage et formées à la forme requise.

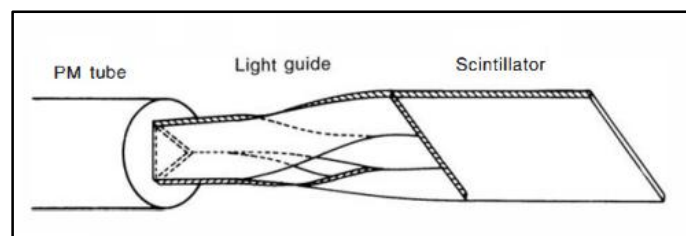


Fig. 05 Guide de lumière à bande utilisé pour coupler le bord d'un grand Scintillateur à un tube PM.

5. Tubes photomultiplicateurs et photodiodes

5.1 Introduction

L'utilisation généralisée du comptage de scintillation dans la détection de rayonnement et la spectroscopie serait impossible sans la disponibilité de dispositifs pour convertir la sortie de lumière extrêmement faible d'une impulsion de scintillation en un signal électrique correspondant. Le tube photomultiplicateur (PM) accomplit cette tâche, convertissant les signaux lumineux qui ne consistent généralement pas en plus de quelques centaines de photons en une impulsion de courant utilisable sans ajouter une grande quantité de bruit aléatoire au signal.

La structure simplifiée d'un tube photomultiplicateur typique est illustrée à la figure (06). Une enveloppe extérieure (généralement en verre) sert de limite de pression pour maintenir les conditions de vide à l'intérieur du tube qui sont nécessaires pour que les électrons de faible énergie puissent être accélérés efficacement par des champs électriques internes. Les deux principaux composants à l'intérieur du tube sont une couche photosensible, appelée photocathode, couplée à une structure multiplicatrice d'électrons. La photocathode sert à convertir autant de photons lumineux incidents en électrons de faible énergie. Si la lumière consiste en une impulsion provenant d'un cristal de scintillation, les photoélectrons produits seront également une impulsion de durée similaire. Étant donné que seules quelques centaines de photoélectrons peuvent être impliqués dans une impulsion typique, leur charge est trop faible à ce stade pour servir de signal électrique pratique. La section multiplicatrice d'électrons dans un tube PM fournit une géométrie de collecte efficace pour les photoélectrons et sert d'amplificateur presque idéal pour augmenter considérablement leur nombre. Après amplification à travers la structure multiplicatrice, une impulsion de scintillation typique donnera lieu à 10^7 - 10^{10} électrons, suffisants pour servir de signal de charge pour l'événement de scintillation d'origine. Cette charge est classiquement collectée à l'anode ou l'étage de sortie de la structure multiplicatrice

5.2 Photocathode

Processus de photoémission

La première étape à effectuer par le tube PM est la conversion des photons lumineux incidents en électrons. Ce processus de photoémission se produit en trois étapes séquentielles :

1. Absorption du photon incident et le transfert d'énergie à un électron dans le matériau photoémissif ;
2. Migration de cet électron vers la surface ;
3. Fuite de l'électron de la surface de la photocathode.

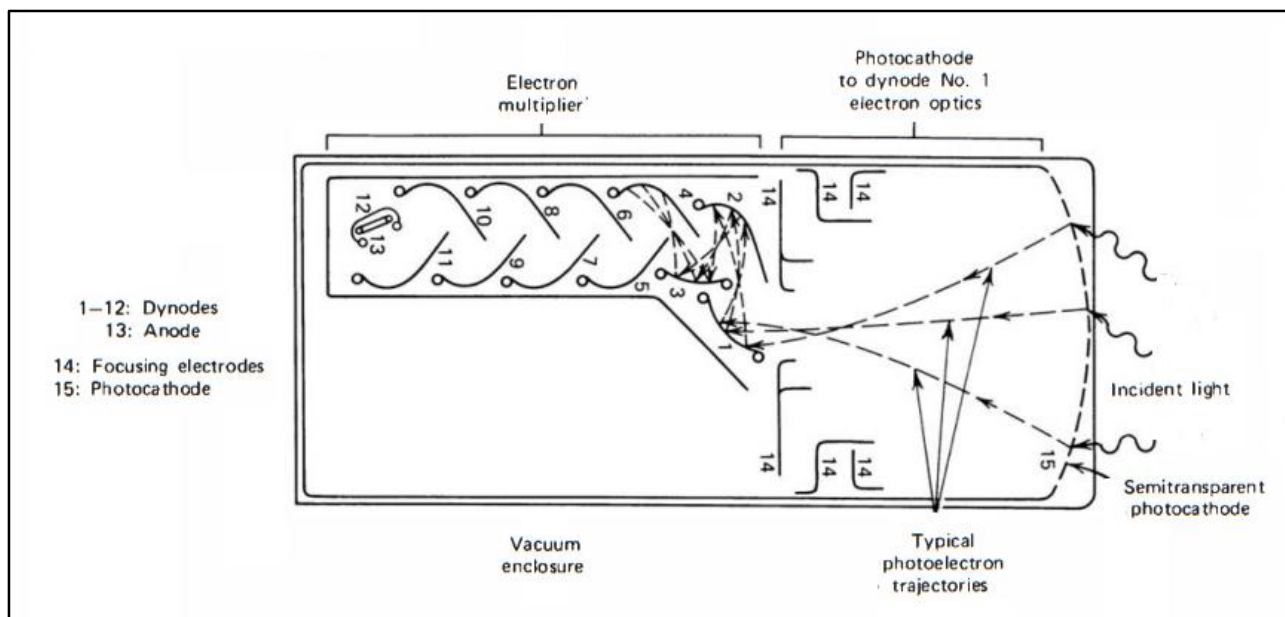


Fig. 06 Éléments de base d'un tube PM

5.3 Multiplication d'électrons

5.3.1 Émission d'électrons secondaires

La partie multiplicateur d'un tube PM est basée sur le phénomène d'électron secondaire. Les électrons de la photocathode sont accélérés et amenés à frapper la surface d'une électrode, appelée dynode. Si le matériau de la dynode est bien choisi, l'énergie déposée par l'électron incident peut entraîner la réémission de plus d'un électron par la même superficie. Le processus d'émission d'électrons secondaires est similaire à celui de la photoémission vue dans la section précédente. Dans ce cas, cependant, les électrons dans le matériau de la dynode sont excités par le passage de l'électron énergétique incident à l'origine sur la surface plutôt que par un photon optique.

Les électrons sortant de la photocathode ont une énergie cinétique de l'ordre de 1 eV ou moins. Ainsi, si la première dynode est maintenue à un potentiel positif de plusieurs centaines de volts, l'énergie cinétique des électrons à l'arrivée à la dynode est déterminée presque entièrement par la grandeur de la tension d'accélération. La création d'un électron excité dans le matériau de la dynode nécessite une énergie au moins égale à la bande interdite, qui peut typiquement être de l'ordre de 2-3 eV.

Par conséquent, il est théoriquement possible qu'un électron incident se crée de l'ordre de 30 électrons excités par 100 V de tension d'accélération. Puisque la direction du mouvement de ces électrons est essentiellement aléatoire, beaucoup d'entre eux n'atteindront pas la surface avant leur désexcitation.

D'autres qui arrivent à la surface auront perdu suffisamment d'énergie pour ne pas pouvoir surmonter la barrière de potentiel en surface et sont donc incapables de s'échapper. Par conséquent, seul une

petite fraction des électrons excités contribue finalement au rendement d'électrons secondaires de la surface de la dynode.

5.3.2 Multiplication par plusieurs étapes

Pour obtenir des gains d'électrons de l'ordre de 10^6 , tous les tubes PM utilisent plusieurs étages. Les électrons sortant de la photocathode sont attirés vers la première dynode et produisent N électrons pour chaque photoélectron incident. Les électrons secondaires qui sont produits à la surface de la première dynode ont à nouveau de très faibles énergies, typiquement quelques eV. Ainsi, ils sont assez facilement guidés par un autre champ électrostatique établi entre la première dynode et une deuxième dynode similaire. Ce processus peut être répété plusieurs fois, avec des électrons secondaires à faible énergie de chaque dynode accélérée vers la dynode suivante.

Une source de tension externe doit être connectée aux tubes PM de manière à ce que la photocathode et chaque étage multiplicateur suivant sont correctement polarisés par rapport à un autre. Parce que les électrons doivent être attirés, la première dynode doit être maintenue à une tension qui est positif par rapport à la photocathode, et chaque dynode suivante doit être maintenue à une tension positive par rapport à la dynode précédente. Pour une collecte efficace des photoélectrons, la tension entre la photocathode et la première dynode est souvent plusieurs fois supérieure à la différence de tension dynode à dynode.

6. Photodiodes comme remplaçant des tubes photomultiplicateurs

6.1 Avantages potentiels

Les tubes photomultiplicateurs sont les amplificateurs de lumière les plus couramment utilisés avec les scintillateurs, à la fois dans le fonctionnement en mode impulsion et en mode courant. Cependant, les progrès dans le développement des semi-conducteurs les photodiodes ont conduit à la substitution de dispositifs à semi-conducteurs pour les tubes PM dans certaines applications. En général, les photodiodes offrent les avantages d'un rendement quantique plus élevé (et donc une meilleure résolution énergétique), une consommation d'énergie plus faible, une taille plus compacte et une robustesse améliorée par rapport aux tubes PM utilisés en scintillation.

Il existe trois conceptions générales qui ont retenu l'attention en tant que substituts possibles des tubes PM. Les photodiodes conventionnelles (souvent appelées photodiodes PIN) n'ont pas de gain interne et fonctionnent en convertissant directement les photons optiques du détecteur à scintillation en paires électron trou qui sont simplement collectées. Les photodiodes à avalanche intègrent un gain interne via l'utilisation de champs électriques plus élevés qui augmentent le nombre de porteurs de charge collectés. La troisième catégorie est un réseau de nombreuses photodiodes à avalanche de petite dimension exploitées au mode Geiger, parfois appelés photomultiplicateurs au silicium.

6.2 Photodiodes conventionnelles

Lorsque la lumière incidence sur un semi-conducteur, des paires électron-trou sont générée d'une manière similaire à celle décrite pour les rayonnements ionisants incidents. Les photons correspondant à la lumière de scintillation typique transportent environ 3-4 eV d'énergie, suffisante pour créer des paires électron-trou dans un semi-conducteur avec une bande interdite d'environ 1-2 eV. La conversion n'est pas limitée par la nécessité que les porteurs de charge s'échappent d'une surface comme dans une photocathode conventionnelle, de sorte que l'efficacité quantique maximale du processus peut être aussi élevée, 60-80%, plusieurs fois plus grand que dans un tube PM. Cette efficacité quantique élevée s'étend également sur une gamme de longueurs d'onde beaucoup plus large que celle des photocathodes dans les tubes PM. Cependant, il n'y a pas amplification interne de cette charge comme dans un tube PM, de sorte que le signal de sortie est plus petit.

Une configuration courante pour une photodiode au silicium est illustrée à la figure (07). La lumière incidente sur une fenêtre d'entrée de la couche P qui est maintenue aussi mince que possible pour améliorer la transmission de la lumière vers le volume actif du silicium. Les électrons et les trous produits par la lumière sont collectés aux limites de la région centrale I entraînée par le champ électrique résultant de la tension appliquée. La charge induite correspondante est traitée dans un préamplificateur attaché pour produire l'impulsion du signal de sortie.

Dans un événement de scintillation typique, seuls quelques milliers de photons visibles sont produits, de sorte que la taille de l'impulsion de charge qui peut être développée est limitée au mieux à pas plus que le même nombre de charges électroniques. En raison de la faible amplitude du signal, le bruit électronique est un problème majeur dans le fonctionnement en mode impulsionnel, en particulier pour les détecteurs de grande surface et les rayonnements à faible énergie.

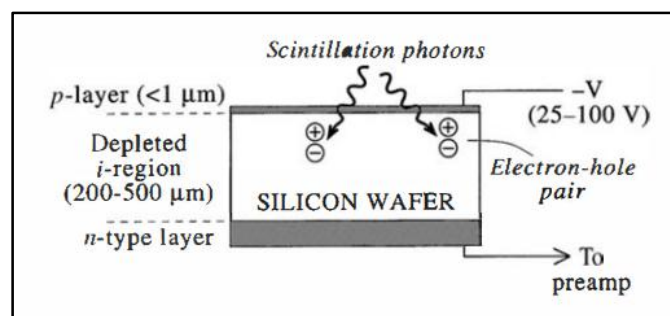


Fig. 07 Configuration de base d'une photodiode conventionnelle

Une approche alternative pour améliorer le facteur de bruit dans les photodiodes consiste à réduire leur capacité. Ces "photodiodes à dérive de silicium" ont des caractéristiques intéressantes puisque leur capacité ne dépend pas directement de leur surface, comme c'est le cas dans d'autres photodiodes. Leur capacité peut être maintenue très faible en incorporant un étage préamplificateur

JFET sur une puce pour éviter la capacité parasite des connexions à un préampli externe. Pour réduire davantage leur niveau de bruit, leur courant de fuite dépendant de la surface peut être réduit par un refroidissement modeste à environ 0 °C. Une petite photodiode à dérivation couplée à un cristal de LaBr₃ de 5 mm de diamètre a démontré une résolution en énergie à température ambiante de 2,7 % à 662 keV, correspondant essentiellement aux meilleures performances des tubes photomultiplicateurs. Encore une fois, l'efficacité quantique plus élevée du silicium est la clé de cette excellente performance. Alors que les photodiodes à dérivation sont limitées à des tailles de cellule unique qui ne dépassent généralement pas 5 à 10 mm de côté, des réseaux monolithiques jusqu'à 77 cellules ont été développés pour fournir une zone sensible de plus de 6 cm² sur une seule puce.

Les photodiodes conventionnelles et à avalanche répondent directement au rayonnement ionisant d'une manière similaire à celle des détecteurs à jonction de silicium. Lorsqu'elle est utilisée pour lire la lumière des scintillateurs, il est possible qu'en plus d'enregistrer les impulsions lumineuses, la photodiode puisse également répondre directement si le rayonnement incident peut pénétrer dans le scintillateur et atteindre le volume actif de la photodiode. Ce contre-effet dit nucléaire peut être particulièrement gênant car la taille de l'impulsion produite par ces interactions directes dans le silicium sera toujours plusieurs fois supérieure aux impulsions produites à partir de la lumière du scintillateur. Les interactions directes dans le silicium produisent une paire électron-trou pour chaque 3,6 e V déposé, alors que 15 à 20 fois cette énergie de dépôt est nécessaire dans un bon scintillateur pour produire la même charge. La différence vient du fait que les scintillateurs ne convertissent généralement qu'environ 5 à 10% de l'énergie des particules en lumière, que toute cette lumière n'entre pas dans la photodiode et qu'il y a moins de 100% d'efficacité quantique pour convertir cette lumière en paires électron-trou. Afin de minimiser cet effet de contre-nucléaire, la photodiode doit être maintenue mince. Cependant, la capacité de la diode augmente à mesure que son épaisseur diminue, ce qui augmente le niveau de bruit. Ainsi, un compromis doit être trouvé entre éviter ces interactions directes et la préservation d'une résolution en énergie raisonnable.

6.3 Photodiodes à avalanche

6.3.1 Propriétés générales

La faible quantité de charge qui est produite dans une photodiode conventionnelle par un événement de scintillation typique peut être augmentée par un processus d'avalanche qui se produit dans un semi-conducteur à des valeurs élevées de la tension appliquée. Les porteurs de charge sont suffisamment accélérés entre les collisions pour créer des paires électron-trou supplémentaires le long du chemin de collecte, en grande partie de la même manière que la multiplication des gaz se produit dans un compteur proportionnel. (Ce même procédé est décrit pour des détecteurs à diodes à avalanche pour des rayonnements ionisants). Le gain interne aide à extraire le signal du niveau de

bruit électronique et permet une bonne résolution énergétique en mode impulsionnel à une énergie de rayonnement inférieure à celle possible avec les photodiodes conventionnelles. Étant donné que le facteur de gain est très sensible à la température et à la tension appliquée, les photodiodes à avalanche nécessitent des alimentations haute tension bien régulées pour un fonctionnement stable. La dépendance du gain à la température est forte par rapport à celle des tubes photo-multiplicateurs, ce qui équivaut à une diminution du gain d'environ 2% par augmentation d'un °C. Pour maintenir une bonne résolution énergétique lorsqu'ils sont utilisés avec des scintillateurs dans des conditions dans lesquelles la température peut varier, des schémas de stabilisation ont été développés sur la base de l'utilisation d'un moniteur de température et de la compensation des ajustements de gain en faisant varier la tension appliquée ou par d'autres moyens électroniques. Pour les applications en mode courant, la stabilité inhérente des photodiodes conventionnelles (sans gain) est généralement préférée.

Les photodiodes à avalanche (PDAs) peuvent être fabriquées de différentes manières, mais un choix courant connu sous le nom de « reach-through » est illustré à la figure (08). La lumière pénètre à travers la fine couche p^+ à gauche du diagramme et interagit quelque part dans la région qui constitue la majeure partie de l'épaisseur de la diode. Les résultats des interactions sont des paires électron-trou, et l'électron est attiré vers la droite à travers la partie de dérive et dans la région de multiplication où existe un champ électrique élevé. Des paires électron-trou supplémentaires sont créées, augmentant le signal mesuré. Des facteurs de gain de quelques centaines sont typiques dans des circonstances normales. Cette amélioration du signal est suffisante pour permettre de détecter des niveaux de lumière beaucoup plus faibles, ou des énergies mesurées dans leur utilisation avec des scintillateurs.

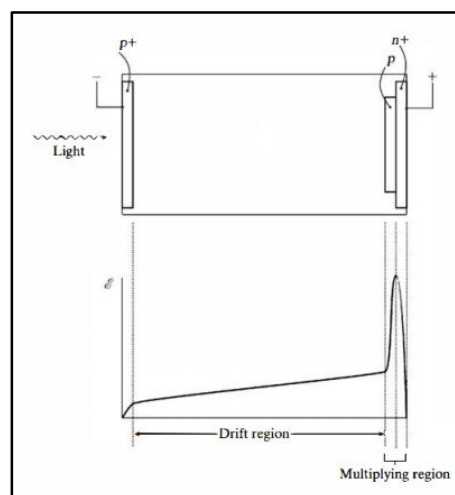


Fig. 08 Configuration de « reach-through » pour une photodiode à avalanche
Champ électrique résultant lorsqu'une tension de polarisation est appliquée

Dans le processus de multiplication, les électrons sont attirés à travers la région de champ élevé et créent des paires électron-trou supplémentaires. Les électrons continuent dans la même direction,

mais les trous seront attirés dans la direction opposée. A des valeurs de champ suffisamment élevées, les trous peuvent également se multiplier et, comme la multiplication des trous produit également des électrons libres supplémentaires, ce processus conduit à un emballement. Cela se produit à la tension de claquage, et la tension appliquée dans des circonstances normales est maintenue quelque peu en dessous de ce niveau. Dans cette région, le gain global sera une fonction exponentielle de la tension appliquée, tenant compte de l'extrême sensibilité à la tension appliquée. Le facteur de gain est également fortement lié à la température, diminuant de quelques pour cent par degré à mesure que la température augmente.

Le processus de multiplication dans une photodiode à avalanche implique que des électrons subissent des collisions à des positions aléatoires tout au long de l'avalanche. Lorsqu'une collision ionisante se produit, un électron libre est augmenté à deux. Les fluctuations de gain d'impulsion à impulsion sont un peu plus importantes que celles observées à partir d'un tube PM typique.

6.3.2 Facteur de bruit excédent

Le comportement statistique du signal de sortie d'un APD ressemble étroitement à celui du compteur proportionnel. Dans les deux cas, les fluctuations d'impulsion à impulsion du signal de sortie résultent de deux effets statistiques indépendants : les variations du nombre de charges initiales (paires électron-trou pour l'APD) pour les dépôts d'énergie fixe dans le scintillateur, et les différences stochastiques dans le processus de multiplication.

Le facteur de bruit en excès J reflète le degré de variabilité au sein de toutes les avalanches qui composent une impulsion. Pour les diodes à avalanche, le facteur de bruit en excès est défini comme :

$$J = 1 + \left(\frac{\sigma_A}{\bar{A}} \right)^2 \quad (09)$$

Où $\bar{A}$ représente l'intensité moyenne de l'avalanche, et σ_A^2 est la variance qui caractérise les fluctuations statistiques d'avalanche à avalanche à électron unique.

7. Photomultiplicateur au silicium

Dans les photodiodes à avalanche, le gain de la diode augmente fortement avec l'augmentation de la tension appliquée. Comme dans les détecteurs à gaz, un niveau est atteint si la tension est suffisamment élevée pour que le processus d'avalanche "s'échappe". Comme la tension de claquage est proche, les régions de multiplication de divers photoélectrons commencent à fusionner pour former une seule avalanche. La diode entre alors dans ce qui peut être le mode dit "Geiger" dans lequel les charges produites lors de l'interaction initiale des photons sont, en principe, multipliés sans limite. Tout comme un tube Geiger rempli de gaz peut produire une décharge d'aussi peu qu'une paire

d'ions, la photodiode à avalanche en mode Geiger peut produire une grande impulsion de sortie à partir d'un seul photon incident.

7.1. Réseau de photodiodes à avalanche en mode Geiger

Pour les applications de scintillation normales dans lesquelles une amplification proportionnelle au nombre original de paires électron-trou est recherchée, ce mode Geiger à cellule unique est de peu d'intérêt puisque toutes les informations sur ce nombre d'origine sont perdues. Cependant, un dispositif utile a été développé qui exploite les propriétés uniques des réseaux de nombreuses photodiodes en mode Geiger de petite dimension.

Le photomultiplicateur en silicium (Si PM) est un réseau de petites cellules de photodiode à avalanche, chacune avec des dimensions de seulement quelques dizaines de microns, produites à l'aide de procédés CMOS sur une puce de silicium. Idéalement, la taille des cellules de photodiode individuelles est suffisamment petite pour qu'il y ait une faible probabilité qu'une cellule donnée soit touchée par un photon de scintillation pendant une impulsion de scintillation, et au plus un seul photon est incident sur une cellule typique. Le nombre de cellules produisant une avalanche est alors proportionnel au nombre de photons de scintillation incidents. Étant donné que les détecteurs à scintillation normaux produiront plusieurs milliers de photons lumineux à leur fenêtre de sortie, le nombre de cellules doit être un grand multiple du nombre de photons collectés pour atteindre cette condition, de sorte que les réseaux de 10^4 cellules ou plus ne sont pas rares. La sortie de chaque cellule lorsqu'elle est utilisée en mode Geiger est très proche de la même amplitude, fixée par l'uniformité des cellules. Puis en ajoutant simplement leur sorties en les connectant en parallèle produit une impulsion analogique dont l'amplitude est proportionnelle au nombre de photons détectés.

7.2 Impulsion et bruits d'obscurité

Bien que l'APD en mode Geiger soit très sensible aux photoélectrons uniques, il est également sensible aux électrons de la bande de conduction générés thermiquement dans le silicium. Ces électrons conduisent à des événements thermiques parasites qui ajoutent une composante de bruit aléatoire au signal de scintillation. Le taux d'obscurité observé à partir d'un Si PM peut atteindre 10^6 impulsions/s par mm^2 à la température ambiante. La plupart de ces événements correspondent à la détérioration d'une seule cellule. Ainsi si la mesure consiste à rechercher des impulsions à partir de photons uniques, le taux d'obscurité spontané sera très élevé. Ce niveau de discrimination serait possible dans de nombreuses applications de la scintillation où un grand nombre de photons sont enregistrés par impulsion. La disponibilité des électrons thermiques peut également être réduite en refroidissant l'appareil. Puisque les réponses de nombreux scintillateurs spectraux les culminent dans

la région bleue où la pénétration de la lumière dans le silicium est très faible, le champ électrique doit être limité très près de la surface pour minimiser les collectes des porteurs.

En plus des événements thermiques, le bruit d'obscurité du Si PM comprend la post-impulsion et la diaphonie optique. La post-impulsion se produit lorsque des impuretés dans le réseau de silicium agissent comme des pièges qui se dés excitent de manière exponentielle dans le temps. Les électrons de ces pièges dés excités peuvent provoquer des avalanches ultérieures dans la même cellule et augmenter le nombre d'avalanches de photons de scintillation. La diaphonie optique se produit en raison de la fluorescence initiée par avalanche dans laquelle les photons émis par la cellule déclenchent des avalanches dans les cellules voisines. Au fur et à mesure que le facteur de remplissage augmente, ces photons d'avalanche sont plus susceptibles de déclencher des avalanches ultérieures. Une solution à ce mécanisme de bruit est la fabrication de tranchées opaques entre pixels. Cependant, lorsqu'un scintillateur enveloppé de manière réfléchissante est accouplé au SiPM, les photons d'avalanche peuvent se refléter à l'intérieur du scintillateur et surmonter toute barrière intercellulaire. Les statistiques du bruit de diaphonie optique dépendront alors également des caractéristiques optiques du scintillateur. Dans l'ensemble, le bruit d'un SiPM est une interaction complexe d'événements de scintillation thermique, post-impulsion et de diaphonie optique.

7.3 Propriétés opérationnelles

L'efficacité de détection de photons (PDE) du SiPM est le produit du facteur de remplissage géométrique, de l'efficacité quantique et de la probabilité de déclenchement d'avalanche. Un facteur de remplissage important nécessite souvent l'utilisation d'une trempe passive, car une seule résistance nécessite moins d'espace en silicium que plusieurs transistors. L'efficacité quantique est définie comme la probabilité qu'un photon entre à la fois dans l'appareil et crée un photoélectron. C'est une fonction forte de la longueur d'onde des photons de scintillation et est liée au coefficient d'absorption optique du substrat semi-conducteur. La dernière composante de la PDE est la probabilité d'initiation d'avalanche et est liée à la probabilité d'extinction pour les processus de ramification. Autrement dit, certaines avalanches peuvent commencer à se former mais ne parviennent pas à atteindre un nombre critique de paires électron-trou nécessaires pour former une décharge complète. Cette perte peut être diminuée en augmentant la polarisation et donc le champ électrique. Actuellement, les SiPM commerciaux sont disponibles avec des PDE bien au-dessus de ceux du tube PM bialkali traditionnel.

Des résolutions temporelles de l'ordre de dizaines de picosecondes sont réalisables pour les photodiodes à avalanche unicellulaires en raison des temps de montée extrêmement rapides et constants qui sont générés par l'avalanche. La synchronisation des réseaux de cellules est souvent de l'ordre de centaines de picosecondes.

Les avantages communément cités du SiPM par rapport au tube PM sont son insensibilité aux champs magnétiques, sa faible tension de fonctionnement et son facteur de forme compact. Des dispositifs à faible coût sont possibles si des conceptions viables peuvent être produites dans des fonderies commerciales de semi-conducteurs. Cependant, plusieurs défis de conception demeurent. La première est que le bruit thermique est proportionnel à la surface du détecteur. Pour une bonne résolution énergétique des grands scintillateurs, la taille SiPM doit être sélectionnée pour maximiser le nombre de photons détectés tout en minimisant le bruit d'obscurité. Un autre défi important est la forte dépendance du gain SiPM, du bruit et de la PDE à la température et à la tension. Pour un fonctionnement stable, des circuits de polarisation fiables doivent être conçus pour mesurer ou compenser toute fluctuation sur la durée de vie de la mesure.

Le SiPM a une réponse intrinsèquement non linéaire aux impulsions de lumière qui sont suffisamment intenses pour que la probabilité ne soit plus faible de déclencher une cellule typique. Supposons un éclair de lumière dont la durée est courte par rapport au temps de récupération d'une cellule illumine uniformément la surface SiPM. Si la probabilité d'activation d'une cellule donnée est de 20 %, alors il y a une probabilité de $(0,2)^2$ ou 4 % que deux photons capables de déclencher la cellule arriveront dans son temps de résolution. La taille de la sortie de la cellule est toujours la même que pour un seul photon, donc certains photons sont "perdus" par ce processus aléatoire et l'impulsion de sortie du SiPM sera inférieure à l'amplitude qui serait produite si chaque photon avait déclenché une cellule individuelle. La perte fractionnelle augmentera avec l'intensité du flash lumineux, de sorte qu'une mesure de l'amplitude de l'impulsion SiPM par rapport à l'intensité lumineuse montrera un écart croissant par rapport à la proportionnalité à mesure que l'intensité lumineuse devient plus grande. En principe, une simple correction basée sur un étalonnage initial peut restaurer la linéarité à condition que les pertes ne représentent pas une grande partie de tous les photons incidents.

8. Analyse de la forme d'impulsion de scintillation

La forme de l'impulsion de tension produite à l'anode d'un tube PM suite à un événement de scintillation dépend de la constante de temps du circuit anodique. Comme il a été indiqué, deux extrêmes peuvent être identifiés, tous deux couramment utilisés en relation avec le comptage par scintillation. La première correspond aux situations dans lesquelles la constante de temps est choisie grande par rapport à la constante de temps de décroissance du scintillateur. C'est la situation généralement choisie si une bonne résolution de hauteur d'impulsion est un objectif majeur et que les fréquences d'impulsion ne sont pas excessivement élevées. Alors chaque impulsion d'électrons est intégrée par le circuit anodique pour produire une impulsion de tension dont l'amplitude est égale à Q/C , le rapport de la charge électronique collectée à la capacité du circuit anodique. Le deuxième cas

est obtenu en ajustant la constante de temps du circuit d'anode pour qu'elle soit beaucoup plus petite que le temps de décroissance du scintillateur. Comme l'analyse suivante le montrera, une impulsion beaucoup plus rapide en résulte, ce qui peut souvent être un avantage dans les applications de temporisation rapide ou lorsque des fréquences d'impulsion élevées sont rencontrées. Dans le même temps, un sacrifice est alors fait dans l'amplitude et la résolution des impulsions.

Le circuit anodique peut être idéalisé comme le montre la figure (09) représente la capacité de l'anode elle-même, plus la capacité du câble de connexion et la capacité d'entrée du circuit auquel l'anode est connectée. La résistance de charge R peut être une résistance physique câblée dans la base du tube ou l'impédance d'entrée du circuit connecté.

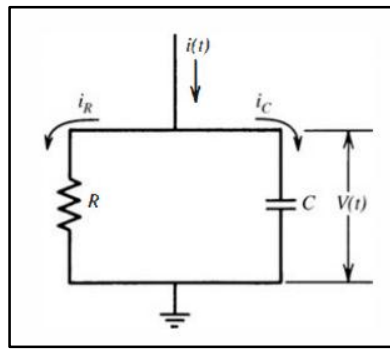


Fig. 09 Circuit RC parallèle représentant le circuit anodique d'un tube PM

Le courant circulant dans l'anode $i(t)$ est simplement le courant d'électrons d'une seule impulsion, supposée commencer à $t = 0$. La forme de $i(t)$ influencera évidemment la forme de l'impulsion de tension d'anode, et nous choisissons pour l'analyse une représentation simplifiée d'une impulsion électronique typique suite à un événement de scintillation. La composante principale de la lumière émise par la plupart des scintillateurs peut être représentée de manière adéquate comme une simple décroissance exponentielle. Si la propagation du temps de transit du tube PM est faible par rapport à ce temps de décroissance, alors un modèle réaliste du courant d'électrons arrivant à l'anode du tube PM est simplement :

$$i(t) = i_0 e^{-\lambda t} \quad (10)$$

Où λ est la constante de désintégration du scintillateur. Le courant initial i_0 peut être exprimé en fonction de la charge totale Q collectée sur toute l'impulsion en notant :

$$Q = \int_0^{\infty} i(t) dt = i_0 \int_0^{\infty} e^{-\lambda t} dt = \frac{i_0}{\lambda} \quad (11)$$

Donc :

$$i_0 = \lambda Q \quad (12)$$

et :

$$i(t) = \lambda Q e^{-\lambda t} \quad (13)$$

Pour dériver l'impulsion de tension $V(t)$ attendue à l'anode, notons d'abord que le courant circulant dans le circuit RC parallèle doit être la somme du courant circulant dans la capacité i_C et du courant traversant la résistance i_R :

$$i(t) = i_C + i_R \quad (14)$$

$$i(t) = C \frac{dV(t)}{dt} + \frac{V(t)}{R} \quad (15)$$

Insertion de l'équation (13) pour $i(t)$ et en divisant par C , on obtient :

$$\frac{dV(t)}{dt} + \frac{1}{RC} V(t) = \frac{\lambda Q}{C} e^{-\lambda t} \quad (16)$$

La solution de cette équation différentielle inhomogène du premier ordre avec la condition initiale $V(0) = 0$ peut être démontrée comme étant :

$$V(t) = \frac{1}{\lambda - \theta} \cdot \frac{\lambda Q}{C} (e^{-\theta t} - e^{-\lambda t}) \quad (17)$$

où $\theta \equiv 1/(RC)$ est l'inverse de la constante de temps de l'anode.

Cas 1. Grande constante de temps

Si la constante de temps de l'anode est rendue grande par rapport au temps de décroissance du scintillateur, alors $\theta \ll \lambda$ et l'équation (17) peut être approximée par :

$$V(t) \cong \frac{Q}{C} (e^{-\theta t} - e^{-\lambda t}) \quad (18)$$

Un tracé de cette forme d'impulsion est illustré à la figure (10). Puisque $\theta \ll \lambda$, la première exponentielle de l'équation (18) décroît lentement et le comportement à court terme est d'environ :

$$V(t) \cong \frac{Q}{C} (1 - e^{-\lambda t}), \quad \left(t \ll \frac{1}{\theta} \right) \quad (19)$$

Les observations importantes suivantes peuvent maintenant être faites :

1. Le front avant de l'impulsion a le comportement temporel $(1 - e^{-\lambda t})$ et son temps de montée est donc déterminé par la constante de décroissance du scintillateur λ . Les scintillateurs rapides ont de grandes valeurs λ qui conduisent à des impulsions à montée rapide.
2. Le front arrière de l'impulsion a le comportement temporel $e^{-\theta t}$ et se désintègre donc à une vitesse déterminée par la constante de temps du circuit anodique $RC \equiv 1/\theta$.
3. L'amplitude de l'impulsion est donnée simplement par Q/C , mais cette valeur n'est atteinte que si $\theta \ll \lambda$. Retraité, la constante de temps du circuit anodique doit être grande devant le temps de décroissance du scintillateur.

La plupart des comptages de scintillation sont effectués dans ce mode car la hauteur d'impulsion est maximisée et les sources de bruit suivantes auront un effet de dégradation minimal sur la résolution de hauteur d'impulsion. De plus, l'amplitude d'impulsion obtenue n'est pas sensible aux changements de résistance de charge ou à de petits changements dans les caractéristiques temporelles de l'impulsion électronique.

L'utilisateur doit alors choisir une constante de temps qui est au moins 5 à 10 fois supérieure au temps de décroissance du scintillateur mais qui n'est pas excessivement longue pour éviter l'accumulation inutile d'impulsions avec le front arrière d'une impulsion précédente à des taux élevés. La constante de temps est déterminée par le produit RC , et la capacité d'anode ou bien la résistance de charge peut être modifiée pour modifier la valeur de RC . Cependant, dans la plupart des applications, c'est la résistance qui doit être adaptée pour atteindre la constante de temps souhaitée, car la capacité est intentionnellement maintenue à sa valeur minimale pour maximiser l'amplitude d'impulsion (Q/C).

Cas 2. Faible constante de temps

Dans le cas extrême opposé, la constante de temps de l'anode est fixée à une faible valeur par rapport au temps de décroissance du scintillateur, ou $\theta \gg \lambda$. Maintenant l'équation (17) devient :

$$V(t) \approx \frac{\lambda}{\theta} \cdot \frac{Q}{C} (e^{-\lambda t} - e^{-\theta t})$$

(20)

Cette forme d'impulsion est également représentée graphiquement sur la Fig. 10. Le comportement aux petites valeurs de t est maintenant :

$$V(t) = \frac{\lambda}{\theta} \cdot \frac{Q}{C} (1 - e^{-\theta t}), \quad \left(t \ll \frac{1}{\lambda} \right)$$

(21)

alors que pour les grandes valeurs de t :

$$V(t) = \frac{\lambda}{\theta} \cdot \frac{Q}{C} e^{-\lambda t}, \quad \left(t \gg \frac{1}{\theta} \right)$$

(22)

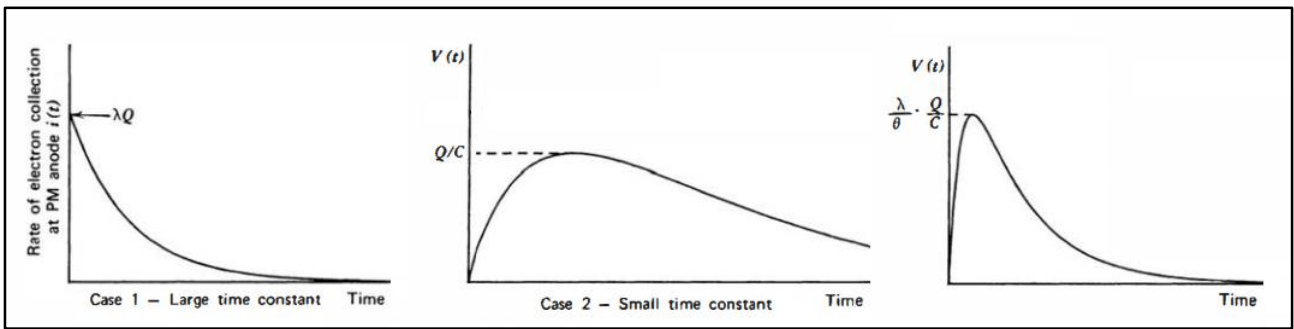


Fig. 10 Pour l'impulsion lumineuse exponentielle supposée représentée à gauche, des tracés sont donnés de l'impulsion d'anode $V(t)$ pour les deux extrêmes de la grande et de la faible constante de temps de l'anode. La durée de l'impulsion est plus courte pour le cas 2, mais l'amplitude maximale est beaucoup plus petite.

Les conclusions générales suivantes s'appliquent désormais :

1. Le front montant de l'impulsion a le comportement temporel $(1 - e^{-\theta t})$, qui est déterminé par la constante de temps d'anode $RC \equiv 1/\theta$.
2. La queue de l'impulsion a le comportement temporel $e^{-\lambda t}$, qui est identique à celui de la lumière du scintillateur.
3. L'amplitude maximale de l'impulsion est maintenant $(\lambda Q/\theta C)$, beaucoup plus petite que le maximum du cas 1 (Q/C) car, par définition du cas 2, $\lambda \ll \theta$.

9. Tubes photomultiplicateurs hybrides

Une variante intéressante de la conception traditionnelle du tube photomultiplicateur (PMT) est le plus souvent appelée tube photomultiplicateur hybride (HPMI) ou, alternativement, photodiode hybride (HPD). Comme le montre la Fig. 11, le principe de base implique une manière fondamentalement différente de multiplier les charges créées dans une photocathode par la lumière incidente. Comme dans un PMT conventionnel, la lumière est convertie en électrons avec une

efficacité quantique dépendant de la longueur d'onde dans une photocathode. La structure classique du multiplicateur d'électrons est maintenant remplacée par un détecteur au silicium placé dans la même enceinte à vide.

Une grande tension, généralement entre 10 et 15 kV, est appliquée entre la photocathode et le détecteur au silicium pour accélérer les électrons à travers le vide entre les deux éléments du tube. Les photoélectrons émergent de la photocathode avec très peu d'énergie, typiquement 1 eV. Cependant, lorsqu'ils sont attirés vers le détecteur au silicium, ils subissent une accélération continue à travers le vide et frappent la surface avant du détecteur au silicium sous forme d'électrons à haute énergie. Par exemple, si une tension de 10 kV est appliquée, alors les électrons arrivent avec une énergie cinétique de 10 keV. Un électron énergétique de ce type perdra son énergie dans le détecteur au silicium par la création de multiples paires électron-trou.

Un électron de 10 keV créera environ 2800 de ces paires si toute son énergie est déposée dans le volume actif du détecteur. Puisque chaque paire porte une charge électronique, le processus qui vient d'être décrit aboutit effectivement à la multiplication de la charge unitaire de chaque photoélectron par un facteur de 2800. Bien que ce niveau d'amplification soit bien inférieur à celui typique des PMT (10^6 à 10^7), il peut néanmoins produire des signaux de taille suffisante pour être amplifiés avec succès dans les composants électroniques suivants.

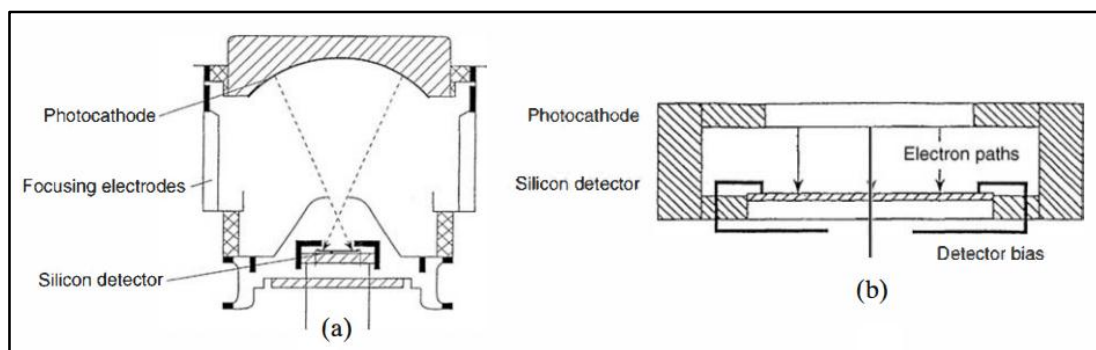


Fig. 11 Eléments de tubes photomultiplicateurs hybrides. La partie (a) montre l'agencement de focalisation électrostatique, tandis que la partie (b) montre une configuration focalisée de proximité

10. Détecteurs à photo ionisation

Il existe une autre alternative aux tubes PM qui a été exploitée pour des applications dans lesquelles la lumière à détecter se situe dans la partie ultraviolette du spectre. Certains composés organiques en phase gazeuse peuvent être ionisés par des photons UV pour former des paires d'ions. Si la vapeur organique est incorporée en tant que composant du gaz de remplissage d'un détecteur conventionnel sensible à l'ionisation (comme un compteur proportionnel), l'amplitude de l'impulsion du signal reflétera le nombre de photons entrants qui ont subi une conversion en ions. De plus, si des

détecteurs sensibles à la position tels que des compteurs proportionnels multifilaires sont utilisés, la position spatiale du point de conversion peut également être déterminée.

Annexe III

A - Détecteurs à diodes semi-conductrices

1. Introduction

Dans de nombreuses applications de détection de rayonnement, l'utilisation d'un milieu de détection solide est d'un grand avantage. Pour la mesure des électrons de haute énergie ou des rayons gamma, les dimensions du détecteur peuvent être maintenu beaucoup plus petites que le détecteur rempli de gaz équivalent car les densités solides sont environ 1000 fois supérieures à celles d'un gaz. Les détecteurs à scintillation offrent à leur tour une possibilité de fournir un milieu de détection solide.

L'une des principales limitations des compteurs à scintillation est leur résolution en énergie relativement faible. La chaîne d'événements qui doit avoir lieu dans la conversion de l'énergie de rayonnement incidente en lumière et la génération subséquente d'un signal électrique implique de nombreuses étapes inefficaces. Par conséquent, l'énergie nécessaire pour produire un porteur d'information (un photoélectron) est de l'ordre de 100 eV ou plus, et le nombre de porteurs créés dans une interaction de rayonnement typique n'est généralement pas supérieur à quelques milliers. Les fluctuations statistiques en si petit nombre imposent une limitation inhérente à la résolution en énergie qui peut être obtenue dans le meilleur des cas, et rien ne peut être fait pour améliorer la résolution en énergie au-delà de ce point.

La seule façon de réduire la limite statistique de la résolution en énergie est d'augmenter le nombre de porteurs d'informations par impulsion. Comme nous le montrons dans ce chapitre, l'utilisation de matériaux semi-conducteurs comme détecteurs de rayonnement peut entraîner un nombre de porteurs beaucoup plus grand pour un événement de rayonnement incident donné que ce qui est possible avec tout autre type de détecteur courant. Par conséquent, la meilleure résolution en énergie des spectromètres de rayonnement utilisés en routine est obtenue à l'aide de détecteurs à semi-conducteurs. Les porteurs d'informations fondamentales sont des paires électron-trou créées le long du chemin emprunté par la particule chargée (rayonnement primaire ou particule secondaire) à travers le détecteur. La paire électron-trou est quelque peu analogue à la paire d'ions créée dans les détecteurs remplis de gaz. Leur mouvement dans un champ électrique appliqué génère le signal électrique de base du détecteur.

En plus d'une résolution énergétique supérieure, les détecteurs à semi-conducteurs peuvent également avoir un certain nombre d'autres caractéristiques souhaitables. Parmi ceux-ci figurent une taille compacte, des caractéristiques de synchronisation relativement rapides et une épaisseur efficace qui peut être modifiée pour répondre aux exigences de l'application.

Les inconvénients peuvent inclure la limitation à de petites tailles et la sensibilité relativement élevée de ces dispositifs à la dégradation des performances due aux dommages induits par les rayonnements.

Parmi les matériaux semi-conducteurs disponibles, le silicium prédomine dans les détecteurs à diode utilisés principalement pour la spectroscopie de particules chargées et discutés dans ce chapitre. Le germanium est plus largement utilisé dans les mesures de rayons gamma, tandis que les dispositifs qui utilisent d'autres matériaux semi-conducteurs seront décrits par la suite.

2. Propriétés des semi-conducteurs

2.1 Structure des bandes dans les solides

Le réseau périodique des matériaux cristallins établit des bandes d'énergie autorisées pour les électrons qui existent dans ce solide. L'énergie de tout électron dans le matériau pur doit être confinée à l'une de ces bandes d'énergie, qui peuvent être séparées par des intervalles ou des plages d'énergies interdites. Une représentation simplifiée des bandes d'intérêt dans les isolants ou les semi-conducteurs, est présentée sur la figure (01). La bande inférieure, appelée bande de valence, correspond aux électrons de la couche externe qui sont liés à des sites de réseau spécifiques dans le cristal. Dans le cas du silicium ou du germanium, ce sont des parties de la liaison covalente qui constituent les forces interatomiques au sein du cristal. La bande supérieure est appelée bande de conduction et représente les électrons libres de migrer à travers le cristal. Les électrons de cette bande contribuent à la conductivité électrique du matériau. Les deux bandes sont séparées par la bande interdite, dont la taille détermine si le matériau est classé comme semi-conducteur ou comme isolant. Le nombre d'électrons dans le cristal est juste suffisant pour remplir complètement tous les sites disponibles dans la bande de valence. En l'absence d'excitation thermique, les isolants comme les semi-conducteurs auraient donc une configuration dans laquelle la bande de valence est complètement pleine et la bande de conduction complètement vide. Dans ces circonstances, aucun des deux ne présenterait théoriquement de conductivité électrique.

Dans un métal, la bande d'énergie occupée la plus élevée n'est pas complètement pleine. Par conséquent, les électrons peuvent facilement migrer à travers le matériau car ils n'ont besoin d'atteindre qu'une faible énergie incrémentielle pour être au-dessus des états occupés. Les métaux sont donc toujours caractérisés par une conductivité électrique très élevée. Dans les isolants ou les semi-conducteurs, en revanche, l'électron doit d'abord traverser la bande interdite pour atteindre la bande de conduction et la conductivité est donc inférieure de plusieurs ordres de grandeur. Pour les isolants, la bande interdite est généralement de 5 eV ou plus, alors que pour les semi-conducteurs, la bande interdite est considérablement inférieure.

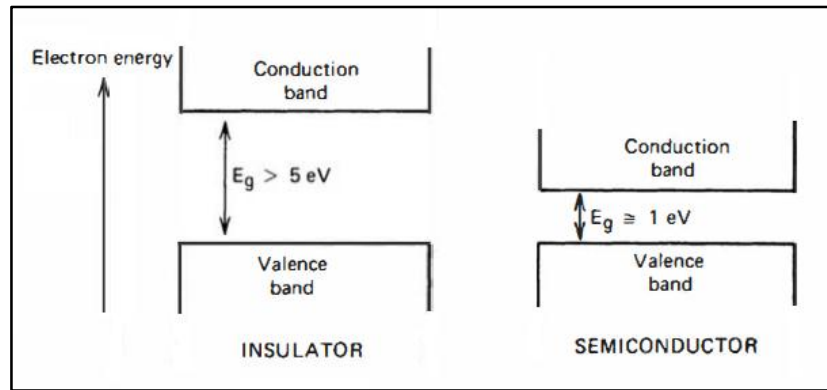


Fig. 01 Structure des bandes d'énergies dans les isolants et les semi-conducteurs

2.2 Porteurs de charge

À toute température non nulle, une certaine énergie thermique est partagée par les électrons du cristal. Il est possible qu'un électron de valence gagne suffisamment d'énergie thermique pour être élevé à travers la bande interdite dans la bande de conduction. Physiquement, ce processus représente simplement l'excitation d'un électron qui fait normalement partie d'une liaison covalente de sorte qu'il peut quitter le site de liaison spécifique et dériver à travers le cristal. Le processus d'excitation crée non seulement un électron dans la bande de conduction autrement vide, mais il laisse également une place vacante (appelée trou) dans la bande de valence autrement pleine. La combinaison des deux s'appelle une paire électron-trou et est à peu près l'analogue à l'état solide de la paire d'ions dans les gaz. L'électron dans la bande de conduction peut être amené à se déplacer sous l'influence d'un champ électrique appliqué. Le trou, représentant une charge positive nette, aura également tendance à se déplacer dans un champ électrique, mais dans une direction opposée à celle de l'électron. Le mouvement de ces deux charges contribue à la conductivité observée du matériau. La probabilité par unité de temps qu'une paire électron-trou soit générée thermiquement est donnée par :

$$p(T) = CT^{3/2} \exp\left(-\frac{E_g}{2kT}\right) \quad (01)$$

Où :

T = température absolue ;

E_g = énergie de bande interdite ;

k = constante de Boltzmann ;

C = constante de proportionnalité caractéristique du matériau.

Comme le reflète le terme exponentiel, la probabilité d'excitation thermique dépend de manière critique du rapport de l'énergie de la bande interdite à la température absolue. Matériaux avec une grande bande interdite aura une faible probabilité d'excitation thermique et par conséquent montrera

une très faible conductivité électrique caractéristique des isolants. Si la bande interdite est aussi faible que quelques électronvolts, une excitation thermique suffisante entraînera une conductivité suffisamment élevée pour que le matériau soit classé comme semi-conducteur. En l'absence d'un champ électrique appliqué, les paires électron-trou créées thermiquement se recombinent finalement, et un équilibre est établi dans lequel la concentration de paires électron-trou observée à un instant donné est proportionnelle au taux de formation. De l'équation (01), cette concentration à l'équilibre est fortement fonction de la température et diminuera drastiquement si le matériau est refroidi.

Après leur formation, l'électron et le trou participent à un mouvement thermique aléatoire qui se traduit par leur diffusion loin de leur point d'origine. Si tous les électrons (ou trous) étaient initialement créés en un seul point, cette diffusion conduit à un élargissement de la répartition des charges en fonction du temps. Une coupe transversale à travers cette distribution serait approximée par une fonction gaussienne avec un écart type a donné par :

$$\sigma = \sqrt{2Dt} \quad (02)$$

où D est le coefficient de diffusion et t est le temps écoulé. Les valeurs de D peuvent être prédites à partir de la relation :

$$D = \mu \frac{kT}{e} \quad (03)$$

où μ est la mobilité du porteur de charge, k est la constante de Boltzmann et T est la température absolue. A 20°C (293 K), la valeur numérique de kT/e est de 0,0253 V.

2.3 Migration des porteurs de charge dans un champ électrique

Si un champ électrique est appliqué au matériau semi-conducteur, les électrons et les trous subiront une migration nette. Le mouvement sera la combinaison d'une vitesse thermique aléatoire et d'une vitesse de dérive nette parallèle à la direction du champ appliqué. Le mouvement des électrons de conduction est un processus relativement facile à visualiser, mais le fait que les trous contribuent également à la conductivité est moins évident. Un trou se déplace d'une position à une autre si un électron quitte un site de valence normal pour remplir un trou existant. La place vacante laissée par l'électron représente alors la nouvelle position du trou. Parce que les électrons seront toujours attirés préférentiellement dans une direction opposée au vecteur champ électrique, les trous se déplacent dans la même direction que le champ électrique. Ce comportement est cohérent avec celui attendu d'une charge positive ponctuelle, car le trou représente en fait l'absence d'électron chargé négativement.

Aux valeurs faibles à modérées de l'intensité du champ électrique, la vitesse de dérive est proportionnelle au champ appliqué. Alors une mobilité μ pour les électrons et les trous peut être définie par :

$$\boxed{v_h = \mu_h \mathcal{E} \quad v_e = \mu_e \mathcal{E}} \quad (04)$$

où $\mathcal{E}$ est l'amplitude du champ électrique. Dans les gaz, la mobilité de l'électron libre est beaucoup plus grande que celle de l'ion positif. Cependant, dans le silicium et le germanium, les deux matériaux semi-conducteurs dominants pour les applications de détection, la mobilité de l'électron et du trou est à peu près du même ordre. Les valeurs numériques des mobilités dans ces matériaux sont données dans le tableau (01).

Tableau : 01 Propriétés du silicium et du germanium intrinsèques

	Si	Ge
Atomic number	14	32
Atomic weight	28.09	72.60
Stable isotope mass numbers	28-29-30	70-72-73-74-76
Density (300 K); g/cm ³	2.33	5.32
Atoms/cm ³	4.96×10^{22}	4.41×10^{22}
Dielectric constant (relative to vacuum)	12	16
Forbidden energy gap (300 K); eV	1.115	0.665
Forbidden energy gap (0 K); eV	1.165	0.746
Intrinsic carrier density (300 K); cm ⁻³	1.5×10^{10}	2.4×10^{13}
Intrinsic resistivity (300 K); $\Omega \cdot \text{cm}$	2.3×10^5	47
Electron mobility (300 K); cm ² /V · s	1350	3900
Hole mobility (300 K); cm ² /V · s	480	1900
Electron mobility (77 K); cm ² /V · s	2.1×10^4	3.6×10^4
Hole mobility (77 K); cm ² /V · s	1.1×10^4	4.2×10^4
Energy per electron-hole pair (300 K); eV	3.62	
Energy per electron-hole pair (77 K); eV	3.76	2.96

À des valeurs du champ électrique élevées, la vitesse de dérive augmente lentement avec le champ.

Finalement, une vitesse de saturation est atteinte devenant indépendante des augmentations supplémentaires du champ électrique. De nombreux détecteurs à semi-conducteur fonctionnent avec

des valeurs de champ électrique suffisamment élevées pour entraîner une vitesse de dérive saturée pour les porteurs de charge. Du fait que ces vitesses saturées sont de l'ordre de 10^7 cm/s, le temps nécessaire pour collecter les porteurs sur des dimensions typiques de 0,1 cm ou moins sera inférieur à 10 ns. Les détecteurs à semi-conducteurs peuvent donc être parmi les plus rapides de tous les types de détecteurs de rayonnement.

En plus de leur dérive, les porteurs de charge subiront également l'influence de la diffusion mentionnée dans la section précédente. Sans diffusion, tous les porteurs de charge se déplaceraient vers les électrodes collectrices en suivant exactement les lignes de champ électrique qui relient leur point d'origine à leur point de collecte. L'effet de la diffusion est d'introduire une certaine dispersion dans la position d'arrivée qui peut être caractérisée comme une distribution gaussienne dont l'écart type peut être prédit en combinant les équations (02), (03) et (04) :

$$\sigma = \sqrt{\frac{2kTx}{e\mathcal{E}}} \quad (05)$$

où x représente la distance de dérive. Dans les détecteurs de faible volume, une valeur typique pour σ serait inférieure à 100 μm . Cet élargissement de diffusion de la distribution de charge limite la précision avec laquelle les mesures de position peuvent être effectuées en utilisant l'emplacement auquel les charges sont collectées au niveau des électrodes dans les détecteurs à semi-conducteurs.

Le temps de collecte des charges est également étalé par diffusion d'une quantité que l'on peut estimer à partir de l'élargissement spatial divisé par la vitesse de dérive, soit généralement inférieure à 1 ns pour les petits volumes. Dans de nombreux cas, ces effets de diffusion sont négligeables et les charges peuvent être représentées comme se déplaçant le long des lignes de champ électrique, toutes avec la même vitesse de dérive. Cependant, les conséquences de la diffusion peuvent devenir importantes pour les détecteurs de grand volume ou lorsqu'il s'agit de mesures de position ou de synchronisation de haute précision.

2.4 Effet des impuretés ou des dopants

2.4.1 Semi-conducteurs intrinsèques

Dans un semi-conducteur complètement pur, tous les électrons de la bande de conduction et tous les trous de la bande de valence seraient créés par une excitation thermique (en l'absence de rayonnement ionisant). Parce que dans ces conditions chaque électron doit laisser un trou derrière lui, le nombre d'électrons dans la bande de conduction doit être exactement égal au nombre de trous dans la bande de valence.

Un tel matériau est appelé semi-conducteur intrinsèque. Ses propriétés peuvent être décrites théoriquement, mais en pratique, ils sont pratiquement impossibles à réaliser. Les propriétés électriques des matériaux réels ont tendance à être dominées par les très faibles niveaux d'impuretés résiduelles ; cela est vrai même pour le silicium et le germanium, qui sont les semi-conducteurs disponibles dans des puretés pratiques les plus élevées.

Dans les discussions qui suivent, nous laisserons n représenter la concentration (nombre par unité de volume) d'électrons dans la bande de conduction. De plus, p représente la concentration de trous dans la bande de valence. Dans le matériau intrinsèque (indice i), l'équilibre établi par l'excitation thermique des électrons de la bande de valence à la conduction et leur recombinaison ultérieure conduit à un nombre égal d'électrons et de trous, où :

$$\boxed{n_i = p_i} \quad (06)$$

Les grandeurs n_i et p_i sont appelées densités de porteurs intrinsèques. De l'équation (01), il est clair que ces densités seront les plus faibles pour les matériaux à large bande interdite et lorsque le matériau est utilisé à basse température. Les densités intrinsèques de trous ou d'électrons à température ambiante sont de $1,5 \times 10^{10} \text{ cm}^{-3}$ dans le silicium et de $2,4 \times 10^{13} \text{ cm}^{-3}$ dans le germanium.

Dans un conducteur métallique, seul le flux d'électrons chargés négativement contribue à sa conductivité électrique. En revanche, le flux d'électrons chargés à la fois négativement et les trous chargés positivement contribuent à la conductivité d'un semi-conducteur intrinsèque. La valeur de la conductivité (ou son inverse, la résistivité ρ) est déterminée par la densité de porteurs intrinsèque n_i et les mobilités μ_h et μ_e des trous et des électrons. Si nous avons une plaque de semi-conducteur d'épaisseur t et de surface A , le courant I qui circulera lorsqu'une tension V est appliquée à travers l'épaisseur est :

$$\boxed{I = \frac{AV}{\rho t} \quad \text{or} \quad \rho = \frac{AV}{It}} \quad (07)$$

Le courant est composé de deux composantes distinctes : le courant dû au flux de trous I_h et celui dû au flux d'électrons I_e . Bien que les deux types de porteurs de charge se déplacent dans des directions opposées, les courants séparés s'additionnent en raison des charges opposées des trous et des électrons. Ainsi, une mesure externe du courant ne peut à elle seule faire la distinction entre le flux de trous ou d'électrons. Le courant total observé sera leur somme, chaque terme étant donné par le produit de la surface, de la densité de porteurs intrinsèque, de la charge électronique e et de la vitesse de dérive du porteur de charge. Ainsi :

$$I = I_e + I_h = An_i e (v_e + v_h) \quad (08)$$

à partir des équations (04) :

$$I = An_i e \mathcal{E} (\mu_e + \mu_h) = An_i e \frac{V}{l} (\mu_e + \mu_h) \quad (09)$$

Combinant :

$$\rho = \frac{1}{en_i (\mu_e + \mu_h)} \quad (10)$$

Insérons les valeurs numériques pour le silicium intrinsèque à température ambiante :

$$\rho = \frac{1}{(1.6 \times 10^{-19} C)(1.5 \times 10^{10} / \text{cm}^3)(1350 + 480) \text{ cm}^2 / V \cdot s}$$

$$\rho = 2.3 \times 10^5 \frac{V \cdot s \cdot \text{cm}}{C} = 230,000 \Omega \cdot \text{cm}$$

Le matériau de silicium actuellement disponible de pureté la plus élevée ne parvient pas à atteindre cette valeur de résistivité théorique en raison de l'effet des impuretés résiduelles.

2.4.2 Semi-conducteurs de type N

Pour illustrer l'effet du dopage sur les propriétés des semi-conducteurs, nous utilisons le silicium cristallin comme exemple. Le germanium et d'autres matériaux semi-conducteurs se comportent de la même manière. Le silicium est tétravalent et dans la structure cristalline normale forme des liaisons covalentes avec les quatre atomes de silicium les plus proches. Un schéma de cette situation est illustré à la figure (02), où chacun des tirets représente un électron de valence normal impliqué dans une liaison covalente. L'excitation thermique dans le matériau intrinsèque consiste à libérer l'un de ces électrons covalents, laissant derrière lui une liaison ou un trou insaturé.

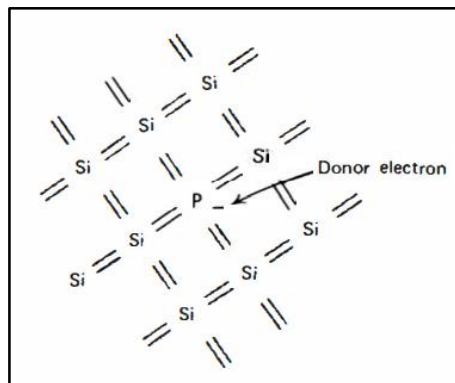


Fig. 02 Représentation d'une impureté donneuse (phosphore) occupant un site de substitution dans un cristal de silicium

Nous considérons maintenant l'effet de la faible concentration d'impuretés qui peuvent être présentes dans le semi-conducteur soit comme une quantité résiduelle après les meilleurs processus de purification, soit comme une petite quantité intentionnellement ajoutée au matériau appelé "dopant" pour adapter ses propriétés. Nous supposons d'abord que l'impureté est pentavalente ou se trouve dans le groupe V du tableau périodique. Lorsqu'il est présent en petites concentrations (de l'ordre de quelques parties par million ou moins), l'atome d'impureté occupera un site de substitution dans le réseau, prenant la place d'un atome de silicium normal. Puisqu'il y a cinq électrons de valence entourant l'atome d'impureté, il en reste un après la formation de toutes les liaisons covalentes. Cet électron supplémentaire est en quelque sorte orphelin et ne reste que très légèrement lié au site d'impureté d'origine. Il faut donc très peu d'énergie pour le déloger pour former un électron de conduction sans trou correspondant. Les impuretés de ce type sont appelées impuretés donneuses car elles apportent facilement des électrons à la bande de conduction. Parce qu'ils ne font pas partie du réseau régulier, les électrons supplémentaires associés aux impuretés donneuses peuvent occuper une position dans l'espace normalement interdit. Ces électrons très faiblement liés auront une énergie proche du sommet de l'écart, comme le montre la figure (03).

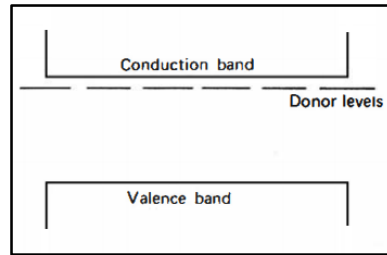


Fig. 03 Niveaux d'énergie des donneurs créés dans la bande interdite du silicium

L'espacement d'énergie entre ces niveaux donneurs et le bas de la bande de conduction est suffisamment petit pour que la probabilité d'excitation thermique donnée par l'équation (01) est suffisamment élevée pour garantir qu'une grande partie des impuretés du donneur sont ionisées. Dans presque tous les cas, la concentration de l'impureté N_D est élevée par rapport à la concentration d'électrons attendue dans la bande de conduction pour le matériau intrinsèque. Par conséquent, le nombre d'électrons de conduction devient complètement dominé par la contribution des impuretés donneuses, et nous pouvons écrire :

$$n \cong N_D \quad (11)$$

La concentration supplémentaire d'électrons dans la bande de conduction par rapport à la valeur intrinsèque augmente le taux de recombinaison, déplaçant l'équilibre entre les électrons et les trous.

En conséquence, la concentration à l'équilibre des trous est diminuée d'une quantité telle que la constante d'équilibre donnée par le produit de n et p est la même que pour le matériau intrinsèque :

$$np = n_i p_i \quad (12)$$

Par exemple, dans le silicium à température ambiante, les densités de porteurs intrinsèques sont d'environ 10^{10} cm^{-3} . Si une impureté donneuse est présente à une concentration de $10^{17} \text{ atomes/cm}^3$ (environ 2 parties par million), la densité d'électrons de conduction n sera de 10^{17} cm^{-3} et la concentration de trous p sera de 10^3 cm^{-3} . Comme le nombre total de porteurs de charge des deux types est maintenant beaucoup plus important (10^{17} cm^{-3} contre $2 \times 10^{10} \text{ cm}^{-3}$), la conductivité électrique d'un semi-conducteur dopé est toujours beaucoup plus grande que celle du matériau pur correspondant.

Même si les électrons de conduction sont maintenant beaucoup plus nombreux que les trous, la neutralité de charge est maintenue en raison de la présence d'impuretés donneuses ionisées. Ces sites représentent des charges positives nettes qui équilibrent exactement les charges électroniques en excès. Ils ne doivent cependant pas être confondus avec des trous car les donneurs ionisés sont fixés dans le réseau et ne peuvent pas migrer.

L'effet net dans le matériau de type N, est donc de créer une situation dans laquelle le nombre d'électrons de conduction est beaucoup plus grand et le nombre de trous beaucoup plus petit que dans le matériau pur. La conductivité électrique est alors déterminée presque exclusivement par le flux d'électrons, et les trous ne jouent qu'un très faible rôle. Dans ce cas, les électrons sont appelés les porteurs majoritaires et les trous les porteurs minoritaires.

La résistivité du matériau dopé peut être calculée à partir de la concentration en dopant et de la mobilité du porteur majoritaire, puisque seul le flux des porteurs majoritaires est important dans les courants mesurés. A titre d'exemple, supposons que nous ayons du silicium avec une densité de donneurs de $10^{13}/\text{cm}^3$, qui sera également la concentration d'électrons de conduction. La résistivité sera alors :

$$\rho = \frac{1}{e N_D \mu_e} \quad (13)$$

$$\rho = \frac{1}{(1.6 \times 10^{-19} \text{ C})(10^{13}/\text{cm}^3)(1350 \text{ cm}^2/\text{V} \cdot \text{s})} = 463 \Omega \cdot \text{cm}$$

2.4.3 Semi-conducteurs de type P

L'ajout d'une impureté trivalente telle qu'un élément du groupe III du tableau périodique à un réseau de silicium aboutit à une situation esquissée sur la figure (04). Si l'impureté occupe un site de substitution, elle a un électron de valence de moins que les atomes de silicium environnants et, par conséquent, une liaison covalente reste non saturée. Ce vide représente un trou similaire à celui laissé derrière lorsqu'un électron de valence normal est excité dans la bande de conduction, mais ses caractéristiques énergétiques sont légèrement différentes. Si un électron est capturé pour combler cette lacune, il participe à une liaison covalente qui n'est pas identique à la masse du cristal car l'un des deux atomes participants est une impureté trivalente. Un électron remplissant ce trou, bien que toujours lié à un emplacement spécifique, est légèrement moins fermement attaché qu'un électron de valence typique. Par conséquent, ces impuretés accepteuses créent également des sites d'électrons dans la bande interdite d'énergie normalement interdite. Dans ce cas, les niveaux d'accepteurs se situent près du fond de l'écart car leurs propriétés sont assez proches des sites occupés par des électrons de valence normaux.

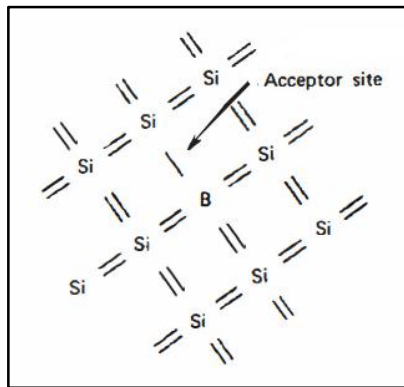


Fig. 04 Représentation d'une impureté acceptrice (bore) occupant
Un site de substitution dans un cristal de silicium

L'excitation thermique normale dans le cristal garantit qu'il y aura toujours des électrons disponibles pour combler les lacunes créées par les impuretés accepteuses ou pour occuper les sites accepteurs illustrés sur la figure (05). Étant donné que la différence d'énergie entre les sites accepteurs typiques et le haut de la bande de valence est faible, une grande partie de tous des sites accepteurs sont remplis par de tels électrons excités thermiquement. Ces électrons proviennent d'autres liaisons covalentes normales dans tout le cristal et laissent donc des trous dans la bande de valence. Pour une bonne approximation, un trou supplémentaire est créé dans la bande de valence pour chaque impureté acceptrice qui est ajoutée. Si la concentration N_A d'impuretés accepteuses est élevée par rapport à la concentration intrinsèque de trous p_i , alors le nombre de trous est complètement dominé par la concentration d'accepteurs, où :

$$p \cong N_A$$

(14)

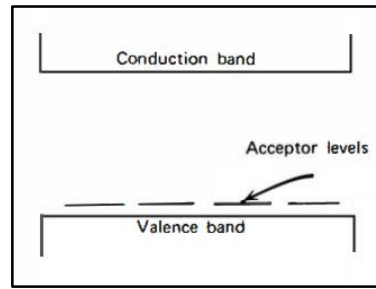


Fig. 05 Niveaux d'énergie des accepteurs créés dans la bande interdite du silicium

La disponibilité accrue de trous améliore la probabilité de recombinaison entre les électrons de conduction et les trous et diminue donc le nombre d'équilibre d'électrons de conduction. Encore une fois, la même constante d'équilibre discutée précédemment est valable, et $n p = n_i p_i$. Dans un matériau de type p, les trous sont les porteurs majoritaires et dominent la conductivité électrique. Les sites accepteurs remplis représentent des charges négatives fixes qui équilibrent la charge positive des trous majoritaires.

2.4.4 Semi-conducteurs compensé

Si des impuretés donneuses et acceptuses sont présentes dans un semi-conducteur en concentration égale, le matériau est dit compensé. Un tel matériau possède certaines des propriétés d'un semi-conducteur intrinsèque car les électrons apportés par les impuretés donneuses sont éliminés dans une certaine mesure par leur capture au site des impuretés acceptuses. Malgré la confusion potentielle avec le matériau intrinsèque purifié, les régions compensées dans les semi-conducteurs reçoivent généralement la désignation i en raison de leurs propriétés quasi intrinsèques.

En pratique, il est impossible d'obtenir une compensation exacte au moment de la fabrication du matériau dopé car tout petit déséquilibre dans la concentration en accepteur ou en donneur conduit rapidement à un comportement de type n ou de type p. Actuellement, le seul moyen pratique d'obtenir une compensation sur de grands volumes de silicium ou de germanium est le processus de dérive des ions lithium après la fabrication du cristal.

2.4.5 Semi-conducteurs fortement dopés

Les couches minces de matériau semi-conducteur qui ont une concentration inhabituellement élevée d'impuretés reçoivent souvent une notation spéciale. Ainsi, n^+ et p^+ désignent des couches fortement dopées de type n et p qui, de ce fait, présentent une conductivité très élevée. Ces couches

sont souvent utilisées pour établir un contact électrique avec des dispositifs semi-conducteurs, car la très faible densité de porteurs minoritaires permet leur application en tant que contacts "bloquants".

2.5 Piégeage et recombinaison

Une fois que les électrons et les trous sont formés dans un semi-conducteur, ils auront tendance à migrer soit spontanément, soit sous l'influence d'un champ électrique appliqué jusqu'à ce qu'ils soient collectés au niveau d'une électrode ou qu'une recombinaison ait lieu. Il existe des prédictions théoriques selon lesquelles la durée de vie moyenne des porteurs de charge avant recombinaison dans des semi-conducteurs parfaitement purs pourrait atteindre une seconde. En pratique, les durées de vie d'au moins trois ou quatre ordres de grandeur inférieures à la seconde effectivement observées sont entièrement dominés par le très faible taux d'impuretés présentes dans le matériau. Certaines de ces impuretés, telles que l'or, le zinc, le cadmium ou d'autres les atomes métalliques occupant des positions de réseau de substitution, introduisent des niveaux d'énergie près du milieu de l'espace interdit. Elles sont donc classées comme impuretés profondes (par opposition aux impuretés acceptrices ou donneuses qui, parce que les niveaux d'énergie correspondants se situent près des bords de la bande interdite, sont appelées impuretés superficielles). Ces impuretés profondes peuvent agir comme des pièges pour les porteurs de charge dans le sens où si un trou ou un électron est capturé, il sera immobilisé pendant une période de temps relativement longue. Bien que le centre de piégeage puisse finalement relâcher la porteuse dans la bande d'où elle provient, le délai est souvent suffisamment long pour empêcher cette porteuse de contribuer à l'impulsion mesurée.

D'autres types d'impuretés profondes peuvent agir comme centres de recombinaison. Ces impuretés sont capables de capturer à la fois les porteurs majoritaires et minoritaires, les obligeant à annihiler. Le niveau d'impureté près du centre de la bande interdit pourrait, par exemple, d'abord capturer une conduction électron. Un peu plus tard, un trou dans la bande de valence pourrait également être capturé, avec l'électron remplit alors le trou. Le site d'impuretés est ainsi ramené à son état d'origine et est capable de provoquer un autre événement de recombinaison. Dans la plupart des cristaux, la recombinaison grâce à de tels centres est bien plus courante que la recombinaison directe d'électrons et de trous sur l'ensemble de la bande interdite.

Le piégeage et la recombinaison contribuent tous deux à la perte de porteurs de charge et tendent à réduire leur durée de vie moyenne dans le cristal. Pour que le matériau serve de bon détecteur de rayonnement, un taux élevé (de préférence 100%) de tous les porteurs créés par le passage du rayonnement incident devraient être collectés. Cette condition restera valable à condition que le temps de collecte des transporteurs soit court par rapport à leur durée de vie moyenne. Les temps de collecte

de l'ordre de 10^{-7} à 10^{-8} s sont assez courant, de sorte que des durées de vie des porteurs de l'ordre de 10^{-5} s ou plus sont généralement suffisantes.

Une autre spécification largement utilisée est la longueur de piégeage dans le matériau. Cette quantité est simplement la distance moyenne parcourue par un porteur avant le piégeage ou la recombinaison et est donné par le produit de la durée de vie moyenne et de la vitesse de dérive moyenne. Dans le but d'avoir un détecteur acceptable, la longueur de piégeage doit être longue par rapport aux dimensions physiques sur lesquelles les porteurs de charges doivent être collectées. Dans les semi-conducteurs à large bande interdite où le piégeage des porteurs charges peut parfois être sévère, il est conventionnel de dire que le produit de la mobilité de la charge et de la durée de vie (appelé « produit mu-tau ») pour quantifier le degré de perte de charges. Ce produit a l'unité de cm^2/V , et lorsqu'il est multiplié par la valeur appliquée l'intensité du champ électrique en V/cm donne la longueur de piégeage.

Outre les impuretés, les défauts structuraux du réseau cristallin peuvent également entraîner au piégeage et à la perte de porteurs de charge. Ces imperfections comprennent des défauts ponctuels tels que des postes vacants ou des interstitiels qui ont tendance à se comporter respectivement comme des accepteurs ou des donneurs. La perte de porteurs de charge peut également se produire au niveau de défauts de ligne ou de dislocations qui peuvent être produites dans des cristaux soumis à des contraintes. Une luxation représente le glissement d'un plan cristallin par rapport à un autre, et son intersection avec la surface du cristal conduit à une piqûre lors de la gravure chimique. La densité de ces piqûres gravées est souvent connue comme mesure de la perfection cristalline d'un échantillon de semi-conducteur.

3. Action des rayonnements ionisants dans les semi-conducteurs

3.1 Energie d'ionisation

Lorsqu'une particule chargée traverse un semi-conducteur avec la structure de bande illustrée dans Fig. (01), l'effet global significatif est la production de nombreuses paires électron-trou le long de la trajectoire de la particule. Le processus de production peut être direct ou indirect, dans la mesure où les particules produisent des électrons de haute énergie (ou rayons delta) qui perdent ensuite leur énergie en produisant plus de paires électron-trou. Quels que soient les mécanismes détaillés impliqués, l'intérêt pratique pour les applications de détection est l'énergie moyenne dépensée par la particule primaire chargée pour produire une paire électron-trou. Cette quantité, souvent appelée l'énergie d'ionisation et étant donné le symbole ϵ , est observée expérimentalement comme étant largement indépendante de l'énergie du rayonnement incident. Cette simplification importante permet d'interpréter le nombre de paires électron-trou produit en termes d'énergie incidente du rayonnement, à condition que la particule soit complètement arrêtée dans le volume actif du détecteur.

Lorsque le rayonnement interagit avec un semi-conducteur, le dépôt d'énergie conduit toujours à la création d'un nombre égal de trous et d'électrons. Cette affirmation est valable indépendamment du fait que le semi-conducteur hôte est pur ou intrinsèque, ou dopé de type p ou de type n. Tout comme des nombres égaux d'électrons libres et d'ions positifs sont créés dans un gaz, chaque électron de conduction produit dans un semi-conducteur doit également créer un trou dans la bande de valence, conduisant à un équilibre du nombre initial de charges créées. Il convient également de souligner que les niveaux de dopage typiques

Dans les semi-conducteurs de type p ou n sont si faibles que ces atomes ne jouent aucun rôle significatif dans la détermination de la nature des interactions des rayonnements dans le matériau. Silicium de type p ainsi que de type n d'épaisseur égale présentera des probabilités d'interaction identiques pour les rayons gamma, et la plage des particules chargées des deux types seront également les mêmes.

L'avantage dominant des détecteurs à semi-conducteurs réside dans la faible énergie d'ionisation. La valeur de E pour le silicium ou le germanium est d'environ 3 eV (voir tableau 01), comparée à avec environ 30 e V requis pour créer une paire d'ions dans les détecteurs à gaz. Ainsi, le nombre de porteurs de charge est 10 fois plus grand pour le boîtier semi-conducteur, pour une énergie donnée déposée dans le détecteur. Le nombre accru de porteurs de charge a deux effets bénéfiques sur la résolution énergétique réalisable. La fluctuation statistique du nombre de porteurs par impulsion devient une fraction plus petite du total à mesure que le nombre augmente. Ce facteur est souvent prédominant dans la détermination de la résolution énergétique limite d'un détecteur pour des environnements moyens à énergie de rayonnement élevée. À basse énergie, la résolution peut être limitée par le bruit électronique dans le préamplificateur, et une plus grande quantité de charge par impulsion conduit à un meilleur rapport signal/bruit.

Un examen plus détaillé montre que ϵ dépend de la nature du rayonnement incident. La plupart des étalonnages des détecteurs sont effectués à l'aide de particules alpha et les valeurs de ϵ indiquées dans le tableau (01) sont basées sur ce mode d'excitation. Toutes les valeurs expérimentales obtenues à l'aide d'autres ions légers et les électrons rapides semblent être assez proches, mais des différences allant jusqu'à 2,2 % ont été signalées entre l'excitation des protons et celle des particules alpha dans le silicium. Ces différences observées montrent la nécessité de réaliser un étalonnage du détecteur avec un type de rayonnement identique à celui impliqué dans la mesure elle-même si des valeurs d'énergie précises sont requises.

Une différence beaucoup plus grande est mesurée pour ϵ lorsque des ions lourds ou des fragments de fission sont impliqués. La valeur de ϵ est nettement plus élevée que pour l'excitation

des particules alpha, ce qui conduit à un nombre de porteurs de charge plus faible que prévu. Les origines physiques de ce défaut de hauteur d'impulsion seront discutées plus loin dans ce chapitre.

L'énergie d'ionisation dépend également de la température. Pour les matériaux de détection les plus importants, la valeur de ϵ augmente avec la diminution de la température. Comme le montre le tableau 01, ϵ dans le silicium est environ 3 % plus élevé à la température de l'azote liquide qu'à la température ambiante.

Il existe également des preuves que l'énergie d'ionisation peut dépendre de l'énergie du rayonnement, en particulier dans la gamme d'énergie des rayons X mous. Il semble que sa valeur augmente avec la diminution de l'énergie des rayons X en dessous d'environ 1 keV, et que le facteur Fano (décrit ci-dessous) augmente également de manière significative avec la diminution de l'énergie dans cette plage.

3.2 Facteur de Fano

Outre le nombre moyen, la fluctuation ou la variation du nombre de porteurs de charge présente également un intérêt primordial en raison du lien étroit entre ce paramètre et la résolution énergétique du détecteur. Comme dans les compteurs à gaz, les fluctuations statistiques observées dans les semi-conducteurs sont plus faibles que prévu si la formation des porteurs de charge était un processus de Poisson. Le modèle de Poisson serait valable si tous les événements sur la trajectoire de la particule ionisante étaient indépendants et prédirait que la variance du nombre total de paires électron-trou devrait être égale au nombre total produit, ou E / ϵ . Le facteur Fano F est introduit comme facteur d'ajustement pour relier la variance observée à la variance prédite de Poisson :

$$F \equiv \frac{\text{observed statistical variance}}{E/\epsilon} \quad (14)$$

Pour une bonne résolution énergétique, on souhaiterait que le facteur de Fano soit le plus petit possible. Bien qu'il n'existe pas encore une compréhension complète de tous les facteurs qui conduisent à une valeur non unitaire pour F , des modèles ont été développés qui rendent compte au moins qualitativement des observations expérimentales. Certaines valeurs numériques pour le silicium et le germanium sont données dans le tableau (01).

Il existe des variations considérables dans les valeurs expérimentales rapportées, en particulier pour le silicium. Le facteur de Fano est généralement mesuré en observant la résolution énergétique d'un détecteur donné dans des conditions dans lesquelles tous les autres facteurs susceptibles d'élargir le pic de pleine énergie (tels que le bruit électronique ou la dérive) peuvent être estimés et pris en

compte. L'hypothèse est alors faite que la largeur résiduelle peut être attribuée uniquement à des effets statistiques. Toutefois, si des facteurs résiduels non statistiques subsistent, le facteur de Fano semblera plus important qu'il ne l'est en réalité. Cela peut expliquer la tendance historique vers des valeurs plus faibles à mesure que les procédures de mesure sont affinées. Il a également été postulé que la valeur du facteur de Fano pourrait dépendre de la nature de la particule qui dépose l'énergie. Certaines mesures suggèrent que sa valeur dans le silicium peut également varier considérablement avec l'énergie du rayonnement, en particulier dans la gamme d'énergie typique des photons X.

4. Semi-conducteurs comme détecteur de rayonnements

4.1 Formation d'impulsions

Lorsqu'une particule dépose de l'énergie dans un détecteur à semi-conducteur, un nombre égal d'électrons de conduction et de trous se forment en quelques picosecondes le long de la trajectoire de la particule. Les configurations de détecteurs abordées dans les sections suivantes garantissent toutes qu'un champ électrique est présent dans tout le volume actif, de sorte que les deux porteurs de charge ressentent des forces électrostatiques qui les font dériver dans des directions opposées. Le mouvement des électrons ou des trous constitue un courant qui persistera jusqu'à ce que ces porteurs soient collectés aux limites du volume actif. Avec l'hypothèse simplificatrice selon laquelle tous les porteurs de charge sont formés en un seul point, les courants résultants peuvent être représentés par le tracé en haut de la figure (06).

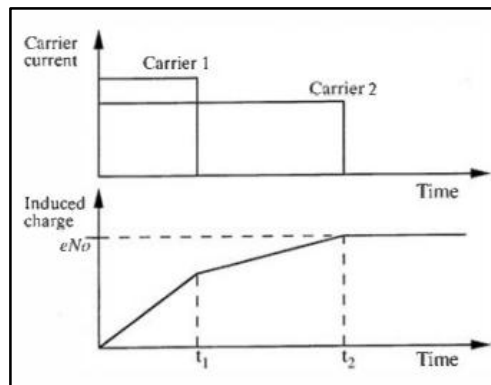


Fig. 06 Le graphique en haut montre une représentation idéalisée des courants d'électrons et de trous circulant dans un semi-conducteur suite à la création de N_0 paires électron-trou.

Dans le graphique du bas, t_1 représente le temps de collecte pour le type de porteur (soit des électrons, soit des trous) qui sont collectés les premiers, et t_2 est le temps de collecte pour l'autre type de porteurs. Si les deux sont entièrement collectés, une charge de eN_0 est induite pour former le signal, où e est la charge de l'électron

Les temps de collecte des charges ne seront probablement pas les mêmes en raison des différences de distance de dérive et les mobilités des porteurs, l'un des deux courants persistera plus

longtemps que l'autre. Lorsque ces courants sont intégrés sur un circuit de mesure avec une grande constante de temps, la charge induite mesurée présente les caractéristiques temporelles indiquées dans le tracé inférieur de la figure 06. Ce profil temporel sera également celui de la montée de l'impulsion produite par un préamplificateur classique utilisé pour traiter les impulsions du détecteur. Ce profil d'impulsion est similaire à celui dérivé pour les chambres ioniques (voir Fig. 11, Annexe I, Détecteurs à gaz, Chambre d'ionisation), à l'exception de l'échelle de temps. Dans les gaz, le temps de collecte des charges positives (ions) est plusieurs fois plus long que celui des charges négatives (électrons libres), de sorte que le mouvement des ions ajoute une très longue composante à l'augmentation de l'impulsion et, d'un point de vue pratique, ne contribue pas à l'impulsion de sortie.

Dans les détecteurs à semi-conducteurs au silicium ou au germanium, la mobilité des trous se situe dans un facteur d'environ 2 ou 3 de la mobilité des électrons, de sorte que les temps de collecte sont beaucoup plus proches d'être équivalents. En conséquence, alors que les chambres ioniques de type impulsionnel n'incluent presque jamais le mouvement des ions dans l'impulsion de sortie, les détecteurs standard à semi-conducteurs au silicium et au germanium reposent sur une intégration complète des courants dus aux électrons et aux trous. Les deux types de porteurs doivent donc être complètement collectés pour que l'impulsion résultante soit une mesure fidèle de l'énergie déposée par la particule.

4.2 Courant de fuite

Afin de créer un champ électrique suffisamment grand pour obtenir une collecte efficace des porteurs de charge à partir de n'importe quel détecteur à semi-conducteur, une tension appliquée généralement de centaines ou de milliers de volts doit être imposée aux bornes du volume actif. Même en l'absence de rayonnement ionisant, tous les détecteurs présenteront une certaine conductivité finie et donc un courant de fuite en régime permanent sera observé. Les fluctuations aléatoires qui se produisent inévitablement dans le courant de fuite auront tendance à obscurcir le petit courant de signal qui circule momentanément après un événement ionisant et représenteront une source de bruit importante dans de nombreuses situations. Les méthodes permettant de réduire le courant de fuite constituent donc une considération importante dans la conception des détecteurs à semi-conducteurs.

La résistivité du silicium de la plus haute pureté actuellement disponible est d'environ 50 000 Ω -cm. Si une plaque de 1 mm d'épaisseur de ce silicium était découpée avec une surface de 1 cm² et équipée de contacts ohmiques, la résistance électrique entre les faces serait de 5 000 Ω . Une tension appliquée de 500 V provoquerait donc un courant de fuite à travers le silicium de 0,1 A. En revanche, le courant de crête généré par une impulsion de 10⁵ porteurs de charge radio-induits ne serait que d'environ 10⁻⁶ A. Il est donc essentiel de réduire ce courant de fuite grâce à l'utilisation de contacts de

blocage. Dans les applications critiques, le courant de fuite ne doit pas dépasser environ 10^{-9} A pour éviter une dégradation significative de la résolution.

À ces niveaux, les fuites à la surface du semi-conducteur peuvent souvent devenir plus importantes que les fuites globales. Un grand soin est apporté à la fabrication des détecteurs à semi-conducteurs pour éviter la contamination des surfaces, qui pourrait créer des chemins de fuite. Certaines configurations peuvent également utiliser des rainures dans la surface ou des anneaux de protection pour aider à supprimer les fuites de surface.

4.3 La jonction semi-conductrice

4.3.1 Propriétés de base des jonctions

Les détecteurs de rayonnement décrits dans ce chapitre sont basés sur les propriétés favorables créées à proximité de la jonction entre les matériaux semi-conducteurs de type n et p. Les porteurs de charge sont capables de migrer à travers la jonction si les régions sont réunies en bon contact thermodynamique. Le simple fait de presser deux morceaux de matériau ne suffira pas, car il restera inévitablement des espaces qui seront grands par rapport à l'espacement du réseau interatomique. En pratique, la jonction est donc normalement formée dans un monocristal en provoquant une modification des conditions de dopage d'un côté à l'autre de la jonction.

À titre d'illustration, supposons que le processus commence avec un cristal de type p dopé avec une concentration uniforme d'impureté accepteur. Dans le profil de concentration en haut de la figure 07, cette concentration initiale d'accepteur N_A est représentée par une ligne horizontale. Nous supposons maintenant que la surface du cristal de gauche est exposée à une vapeur d'impureté de type n qui se diffuse à une certaine distance dans le cristal. Le profil d'impuretés donneuses résultant est étiqueté N_D sur la figure et diminue en fonction de la distance à la surface. Près de la surface, les impuretés donneuses peuvent être amenées à être plus nombreuses que les accepteurs, convertissant ainsi la partie gauche du matériau à un cristal de type n.

La variation approximative de la concentration en porteurs de charge à l'équilibre est également tracée en haut de la figure (07) et étiquetée par p (concentration en trous) et n (concentration en électrons de conduction).

Ces profils sont ensuite modifiés au voisinage de la jonction p-n en raison des effets de diffusion des porteurs de charge. Dans la région de type n à gauche, la densité des électrons de conduction est beaucoup plus élevée que celle du type p. La jonction entre les deux régions représente donc une discontinuité dans la densité électronique de conduction. Partout où un gradient aussi marqué existe pour tout porteur libre de migrer, une diffusion nette des régions de forte concentration vers celles de faible concentration doit avoir lieu.

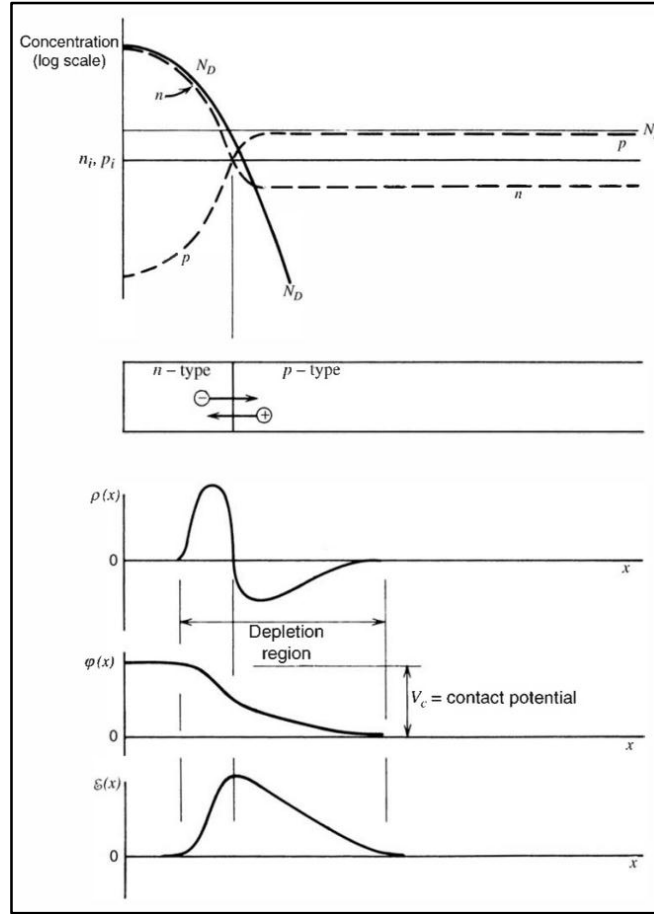


Fig. 07 Les profils de concentration ρ pour la jonction n-p indiqués en haut sont expliqués dans le texte. Les effets de la diffusion des porteurs à travers la jonction donnent lieu aux profils illustrés pour la charge d'espace $\rho(x)$, le potentiel électrique $\phi(x)$ et le champ électrique $\xi(x)$.

Nous supposons maintenant que la surface du cristal de gauche est exposée à une vapeur d'impureté de type n qui se diffuse à une certaine distance dans le cristal. Le profil d'impuretés donneuses résultant est appelé N_D sur la figure et diminue en fonction de la distance à la surface. Près de la surface, les impuretés donneuses peuvent être amenées à être plus nombreuses que les accepteurs, convertissant ainsi la partie gauche du cristal en un matériau de type n.

La variation approximative de la concentration en porteurs de charge à l'équilibre est également tracée en haut de la figure (07) et étiquetée par p (concentration en trous) et n (concentration en électrons de conduction).

Ces profils sont ensuite modifiés au voisinage de la jonction p-n en raison des effets de diffusion des porteurs de charge. Dans la région de type n de gauche, la densité des électrons de conduction est très élevée par rapport au type p. La jonction entre les deux régions présente donc une discontinuité dans la densité électronique de conduction. Partout où un gradient aussi prononcé existe pour tout

porteur libre de migrer, une diffusion nette des régions de forte concentration vers celles de faible concentration doit avoir lieu.

Ainsi, il y aura une certaine diffusion nette des électrons de conduction dans le matériau de type p, où ils se combineront rapidement avec les trous. En effet, cette annihilation représente la capture de l'électron de conduction par l'une des lacunes existant dans les liaisons covalentes du matériau de type p. La diffusion des électrons de conduction hors d'un matériau de type n alors laisse derrière elle des charges positives immobiles sous la forme d'impuretés donneuses ionisées. Un argument similaire et symétrique conduit à la conclusion que les trous (la majorité dans le matériau de type p) doivent également diffuser à travers la jonction car ils voient également un gradient de densité abrupt. Chaque trou retiré du côté p de la jonction laisse derrière lui un site accepteur qui a capté un électron supplémentaire et représente donc une charge négative fixe et immobile. L'effet combiné est de créer une charge d'espace négative nette du côté p et une charge d'espace positive du côté n de la jonction.

La charge d'espace accumulée crée un champ électrique qui diminue la tendance à une diffusion ultérieure. À l'équilibre, le champ est juste suffisant pour empêcher une diffusion nette supplémentaire à travers la jonction, et une distribution de charge en régime permanent est donc établie.

La région dans laquelle existe le déséquilibre de charge est appelée région d'appauvrissement (déplétion) et s'étend à la fois sur les côtés p et n de la jonction. Si les concentrations de donneurs du côté n et d'accepteurs du côté p sont égales, les conditions de diffusion sont approximativement les mêmes pour les trous et les électrons, et la région d'appauvrissement s'étend sur des distances égales dans les deux côtés. Généralement, il existe une différence marquée dans les niveaux de dopage d'un côté de la jonction par rapport à l'autre. Par exemple, si la concentration du donneur dans un matériau de type n est supérieure à celle des atomes accepteurs de type p, les électrons diffusant à travers la jonction auront tendance à parcourir une plus grande distance dans le matériau de type n avant de tous se recombiner avec des trous. Dans ce cas, la région d'appauvrissement s'étendrait plus loin du côté p.

L'accumulation de charge nette dans la région de la jonction conduit à l'établissement d'une différence de potentiel électrique aux bornes de la jonction. La valeur du potentiel ϕ en tout point peut être trouvée par à partir de la solution de l'équation de Poisson :

$$\boxed{\nabla^2 \phi = -\frac{\rho}{\epsilon}} \quad (15)$$

Où ϵ est la constante diélectrique du milieu et ρ est la densité de charge nette. En une dimension, l'équation (15) prend la forme :

$$\boxed{\frac{d^2\varphi}{dx^2} = -\frac{\rho(x)}{\epsilon}} \quad (16)$$

De sorte que la forme du potentiel aux bornes de la jonction peut être obtenue en intégrant deux fois le profil de distribution de charge $\rho(x)$. Des exemples graphiques sont présentés sur la figure (07). A l'équilibre, la différence de potentiel aux bornes de la jonction (appelée potentiel de contact) équivaut à presque la valeur totale de la bande interdite du matériau semi-conducteur. La direction de cette différence de potentiel est telle qu'elle s'oppose à la diffusion ultérieure des électrons de gauche à droite et des trous de droite à gauche sur la figure (07).

Lorsqu'il existe une différence de potentiel électrique, il doit également exister un champ électrique $\mathcal{E}(x)$. Son ampleur se trouve en prenant le gradient du potentiel :

$$\boxed{\mathcal{E} = -\text{grad } \varphi} \quad (17)$$

Ce qui, en une dimension, est simplement :

$$\boxed{\mathcal{E}(x) = -\frac{d\varphi}{dx}} \quad (18)$$

Le champ électrique s'étendra sur toute la largeur de la région de déplétion, dans laquelle le déséquilibre de charge est important et le potentiel présente un certain gradient. Sa variation est également esquissée sur la figure (07).

La région de déplétion présente des propriétés très attractives en tant que milieu pour la détection des radiations. Le champ électrique qui existe fait que tous les électrons créés dans ou à proximité de la jonction doivent être ramenée vers le matériau de type n, et tous les trous sont balayés de la même manière vers le côté de type p. La région est ainsi « appauvrie » dans le sens où la concentration de trous et d'électrons est fortement supprimée. Les seules charges significatives restant dans la région d'épuisement sont les sites donneurs ionisés immobiles et les sites accepteurs remplis. Etant donné que ces dernières charges ne contribuent pas à la conductivité, la région de déplétion présente une résistivité très élevée par rapport aux matériaux de type n et p de chaque côté de la jonction. Les paires électron-trou créées dans la région de déplétion par le passage du rayonnement seront balayées hors de la région de déplétion par le champ électrique et leur mouvement constitue un signal électrique de base.

La génération thermique de porteurs de charge continuera à avoir lieu dans la région de déplétion, apportant une composante parfois appelée courant de génération au courant de fuite. Ces charges sont généralement balayées en quelques nanosecondes, un temps qui est de plusieurs ordres de grandeur plus court que le temps nécessaire pour établir l'équilibre thermique. Ainsi, la

concentration de porteurs à l'état d'équilibre est fortement réduite lors de la région de déplétion car la suppression des charges est un processus beaucoup plus rapide que leur création. La faible concentration de porteurs créée par une particule ionisante est donc facilement détectée au-dessus de cette concentration hautement supprimée et générée thermiquement.

4.3.2 Polarisation inverse

Jusqu'à présent, nous avons parlé d'une jonction de diodes semi-conductrices à laquelle aucune tension externe n'est appliquée. Une telle jonction impartiale fonctionnera comme un détecteur, mais avec de très mauvaises performances. Le potentiel de contact d'environ 1 V qui se forme spontanément aux bornes de la jonction est insuffisant pour générer un champ électrique suffisamment important pour faire bouger très rapidement les porteurs de charge. Par conséquent, les charges peuvent être facilement perdues à la suite du piégeage et de la recombinaison, ce qui entraîne souvent une collecte incomplète des charges. L'épaisseur de la région de déplétion est assez petite et la capacité d'une jonction non polarisée est élevée. Par conséquent, les propriétés de bruit d'une jonction non polarisée connectée à l'étage d'entrée d'un préamplificateur sont assez médiocres. Pour ces raisons, les jonctions non polarisées ne sont pas utilisées comme détecteurs de rayonnement en mode impulsif, mais une tension externe est appliquée dans la direction pour provoquer la polarisation inverse de la diode semi-conductrice.

La jonction (p-n) est plus connue dans son rôle de diode. Les propriétés de la jonction sont telles qu'elle conduira facilement le courant lorsqu'une tension est appliquée dans le sens "vers l'avant", mais elle conduira très peu de courant lorsqu'elle sera polarisée dans le sens "inverse". Dans la configuration de la figure (08), supposons d'abord qu'une tension positive est appliquée au côté p de la jonction par rapport au côté n. Le potentiel aura tendance à attirer les électrons de conduction du côté n ainsi que les trous du côté p à travers la jonction. Comme, dans les deux cas, ce sont les porteurs majoritaires, la conductivité à travers la jonction est grandement améliorée. Le potentiel de contact illustré sur la figure 11.8 est réduit de la quantité de tension de polarisation appliquée, ce qui tend à réduire la différence de potentiel vue par un électron d'un côté à l'autre de la jonction. Il s'agit de la direction de la polarisation directe, et seules de petites valeurs de la tension de polarisation directe sont nécessaires pour que la jonction conduise des courants importants.

Si la situation est inversée et que le côté p de la jonction devient négatif par rapport au côté n, la jonction est polarisée en inverse. Désormais, la différence de potentiel naturelle d'un côté à l'autre de la jonction est renforcée, comme le montre la figure 08 c. Dans ces circonstances, ce sont les porteurs minoritaires (trous du côté n et électrons du côté p) qui sont attirés à travers la jonction et, parce que leur concentration est relativement faible, le courant inverse aux bornes de la diode est

assez faible. Par conséquent, la jonction p-n sert d'élément redresseur, permettant une circulation relativement libre du courant dans un sens tout en présentant une grande résistance à son écoulement dans le sens opposé. Si la polarisation inverse est très importante, une détérioration, soudaine de la diode se produira et le courant inverse augmentera brusquement, souvent avec des effets destructeurs.

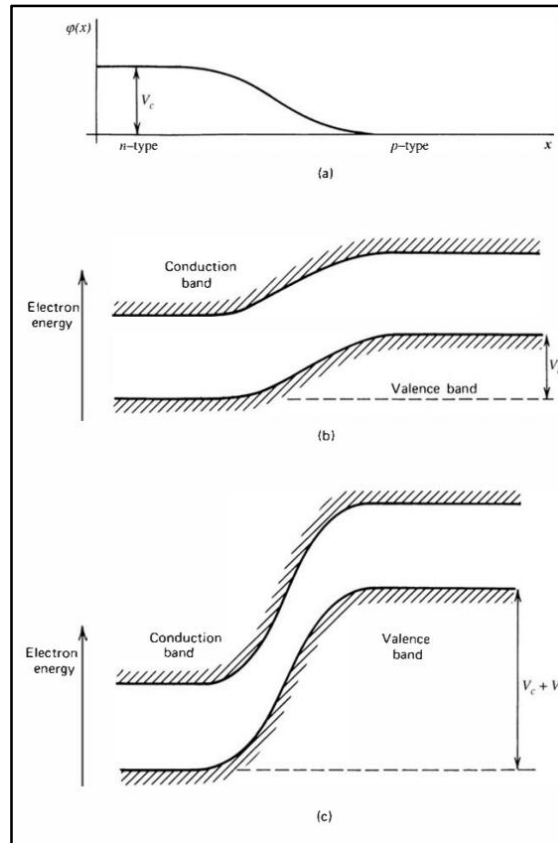


Fig. 08 (a) La variation du potentiel électrique à travers une jonction n-p. La variation résultante des bandes d'énergie électronique à travers la jonction. La courbure est inversée car une augmentation de l'énergie électronique correspond à une diminution du potentiel électrique conventionnel $\phi(x)$ défini pour une charge positive. (c) Le déplacement supplémentaire des bandes provoqué par l'application d'une polarisation inverse V à travers la jonction.

4.3.3 Propriétés de la jonction polarisée en inverse

Lorsqu'une polarisation inverse est appliquée à la jonction, pratiquement toute la tension appliquée apparaîtra dans la région de déplétion, car sa résistivité est beaucoup plus élevée que celle des matériaux type n et p. Parce que l'effet de la polarisation inverse est d'accentuer la différence de potentiel aux bornes de la jonction, l'équation de Poisson (Eq. 15) exige que la charge d'espace augmente également et s'étend sur une plus grande distance de chaque côté de la jonction. Ainsi, l'épaisseur de la région d'appauvrissement augmente également, étendant le volume sur lequel les porteurs de charge produits par le rayonnement seront collectés. Les détecteurs pratiques fonctionnent presque toujours avec une tension de polarisation très élevée par rapport au potentiel de contact, de

sorte que la tension appliquée domine complètement l'amplitude de la différence de potentiel aux bornes de la jonction.

Dans l'analyse qui suit, nous supposons d'abord que la plaquette semi-conductrice dans laquelle la jonction est formée est suffisamment épaisse pour que la région de déplétion n'atteigne aucune des surfaces et soit contenue dans le volume intérieur de la plaquette. Cette condition est valable pour les détecteurs partiellement épuisés dans lesquels une partie de l'épaisseur de la tranche reste intacte. De nombreux détecteurs à semi-conducteurs fonctionnent avec une tension de polarisation inverse suffisante pour que la région de déplétion s'étende sur toute l'épaisseur de la tranche, créant ainsi un détecteur entièrement appauvri (ou totalement appauvri). Ces configurations partagent un bon nombre des propriétés dérivées ci-dessous, sauf que la région de déplétion est évidemment limitée par l'épaisseur physique de la tranche.

Certaines propriétés de la jonction de polarisation inverse peuvent être dérivées si nous représentons la distribution de charges esquissée sur la figure (07) par la distribution idéalisée indiquée ci-dessous :

$$\rho(x) = \begin{cases} eN_D & (-a < x \leq 0) \\ -eN_A & (0 < x \leq b) \end{cases} \quad (19)$$

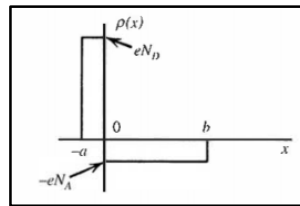


Fig. 09 La distribution de charge esquissée sur la figure (07)

Ici, la diffusion électronique est supposée entraîner une charge d'espace positive uniforme (les sites donneurs ionisés) sur la région $-a < x \leq 0$ du côté n de la jonction. Un négatif correspondant la charge d'espace (les sites accepteurs remplis) résultant de la diffusion des trous est supposée s'étendre sur la région $0 < x \leq b$ du côté p. Parce que la charge nette doit être nulle, $N_D a = N_A b$.

L'équation 16 appliquée à ce cas prend la forme :

$$\frac{d^2\phi}{dx^2} = \begin{cases} -\frac{eN_D}{\epsilon} & (-a < x \leq 0) \\ +\frac{eN_A}{\epsilon} & (0 < x \leq b) \end{cases} \quad (20)$$

Nous effectuons maintenant une intégration et appliquons les conditions aux limites selon lesquelles le champ électrique $\mathcal{E} = -d\phi/dx$ doit disparaître aux deux bords de la distribution de charge :

$$\boxed{\frac{d\varphi}{dx}(-a) = 0 \quad \text{and} \quad \frac{d\varphi}{dx}(b) = 0} \quad (21)$$

Le résultat est alors

$$\boxed{\frac{d\varphi}{dx} = \begin{cases} -\frac{eN_D}{\epsilon}(x+a) & (-a < x \leq 0) \\ +\frac{eN_A}{\epsilon}(x-b) & (0 < x \leq b) \end{cases}} \quad (22)$$

La forme correspondante du champ électrique $\mathcal{E} = -d\varphi/dx$ est esquissée ci-dessous :

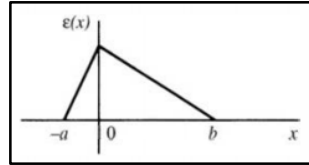


Fig. 10 La forme du champ électrique

Une autre intégration produira maintenant le potentiel électrique $\varphi(x)$. La différence de potentiel du côté n au côté p de la jonction, si l'on néglige le potentiel de contact relativement faible, est juste la valeur de la tension inverse appliqué V . On peut donc appliquer les conditions aux limites :

$$\boxed{\varphi(-a) = V \quad \text{and} \quad \varphi(b) = 0} \quad (23)$$

La solution prend alors la forme :

$$\boxed{\varphi(x) = \begin{cases} -\frac{eN_D}{2\epsilon}(x+a)^2 + V & (-a < x \leq 0) \\ +\frac{eN_A}{2\epsilon}(x-b)^2 & (0 < x \leq b) \end{cases}} \quad (24)$$

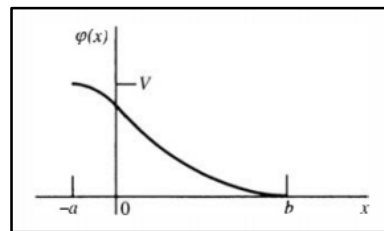


Fig. 11 La forme du potentiel électrique $\varphi(x)$

Puisque les solutions de chaque côté de la jonction doivent correspondre en $x = 0$, nous pouvons écrire :

$$\boxed{V - \frac{eN_D a^2}{2\epsilon} = \frac{eN_A b^2}{2\epsilon} \quad \left| \quad N_A b^2 + N_D a^2 = \frac{2\epsilon V}{e}} \quad (25)$$

Maintenant que $N_D a = N_A b$, l'expression ci-dessus peut être réécrite :

$$(a + b)b = \frac{2\epsilon V}{eN_A} \quad (26)$$

La largeur totale de la région d'appauvrissement d est la distance totale sur laquelle s'étend la charge d'espace, d'où $d = a + b$.

Nous supposons que le niveau de dopage du côté n est beaucoup plus élevé que du côté p , de sorte que $N_D \gg N_A$. Puisque $N_D a = N_A b$, il s'ensuit que $b \gg a$, et donc la charge d'espace s'étend beaucoup plus loin du côté p que du côté n . Alors $d \cong b$ et on peut écrire :

$$d \cong \left(\frac{2\epsilon V}{eN_A} \right)^{1/2} \quad (27)$$

Si nous sommes partis de l'hypothèse inverse selon laquelle le niveau de dopage du côté p était prédominant, un résultat similaire aurait été obtenu sauf que N_A dans l'expression ci-dessus serait remplacé par N_D . Une solution généralisée pour l'épaisseur de la région de déplétion est donc :

$$d \cong \left(\frac{2\epsilon V}{eN} \right)^{1/2} \quad (28)$$

Dans cette expression, N représente maintenant la concentration des dopants (donneurs ou accepteurs) du côté de la jonction qui a le niveau de dopant le plus faible.

La résistivité ρ_d du semi-conducteur dopé (voir Eq. 13) est donné par $1/e\mu N$, où μ est la mobilité du porteur majoritaire. L'équation (28) peut ainsi s'écrire :

$$d \cong (2\epsilon V \mu \rho_d)^{1/2} \quad (29)$$

Etant donné que l'on souhaite souvent la plus grande largeur possible de la zone de déplétion pour une tension donnée, il est avantageux d'avoir la résistivité aussi élevée que possible. Cette résistivité est limitée par la pureté du matériau semi-conducteur avant le processus de dopage, car suffisamment de dopant doit être ajouté pour annuler les effets non uniformes des impuretés

résiduelles. Une priorité est donc accordée à l'obtention de détecteurs fabriqués à partir de matériaux de la plus haute pureté possible.

En raison des charges fixes accumulées de chaque côté de la jonction, la région d'appauvrissement présente certaines propriétés d'un condensateur chargé. Si la polarisation inverse augmente, la région d'appauvrissement devient plus épaisse et la capacité représentée par les charges séparées diminue donc. La valeur de la capacité par unité de surface est :

$$C = \frac{\epsilon}{d} \cong \left(\frac{e\epsilon N}{2V} \right)^{1/2} \quad (30)$$

Une bonne résolution énergétique dans des conditions dans lesquelles le bruit électronique est dominant dépend de l'obtention d'une faible capacité du détecteur et est donc favorisée par l'utilisation de la tension appliquée la plus élevée possible, jusqu'à ce que le détecteur soit complètement épuisé.

Le champ électrique maximal se produira au point de transition entre les matériaux de type n et p. Sa grandeur est donnée par :

$$\mathcal{E}_{\max} \cong \frac{2V}{d} = \left(\frac{2VNe}{\epsilon} \right)^{1/2} \quad (31)$$

et peut facilement atteindre 10^6 - 10^7 V/m dans des conditions typiques. Pour les jonctions partiellement appauvries, l'épaisseur de la couche d'appauvrissement est proportionnelle à $\sqrt{V}$ de sorte que la valeur de $\mathcal{E}_{\max}$ augmente avec la tension appliquée comme $\sqrt{V}$.

La tension de fonctionnement maximale de tout détecteur à diode doit être maintenue en dessous de la tension de claquage pour éviter une détérioration des propriétés du détecteur. Les détecteurs commerciaux sont fournis avec une tension nominale maximale qui doit toujours être strictement respectée. Une protection supplémentaire peut être fournie en surveillant le courant de fuite pendant l'application de la tension.

Pour résumer, la jonction polarisée en inverse constitue un détecteur de rayonnement attrayant car les porteurs de charge créés dans la région d'appauvrissement peuvent être collectés rapidement et efficacement. La largeur de la région d'appauvrissement représente le volume actif du détecteur et est modifiée dans les détecteurs partiellement appauvris en faisant varier la polarisation inverse. Le volume actif variable des jonctions semi-conductrices est unique parmi les détecteurs de rayonnement et est parfois utilisé à bon escient. La capacité d'un détecteur partiellement épuisé varie également en

fonction de la tension appliquée, et un fonctionnement stable nécessite donc l'utilisation de préamplificateurs sensibles à la charge.

5. Configuration des détecteurs de rayonnements à semi-conducteurs

Il existe plusieurs configurations telles que :

- Détecteurs à jonctions diffusées ;
- Détecteurs à barrière de surface ;
- Couches d'ions implantées ;
- Détecteurs entièrement appauvris ;
- Détecteurs planaires passivés ;

6. Caractéristiques opérationnelles

6.1 Courant de fuite

Lorsqu'une tension est appliquée à un détecteur à jonction, c'est-à-dire pour inverser la polarisation de la jonction, un petit courant de l'ordre d'une fraction de microampère est observé. Les origines de ce courant de fuite sont liées à la fois au volume apparent et à la surface du détecteur. Les courants de fuite massifs apparaissant à l'intérieur du volume du détecteur peuvent être provoqués par l'un ou l'autre de deux mécanismes.

6.2 Bruit du détecteur et résolution de l'énergie

Les sources de bruit électronique dans les mesures spectroscopiques se répartissent en deux catégories principales : le bruit série et le bruit parallèle. Pour le type de détecteurs à diode au silicium utilisé pour les mesures de particules chargées, trois principaux contributeurs au bruit électronique sont les plus significatifs :

1. Fluctuations du courant de fuite généré en masse, une composante du bruit parallèle.
2. Fluctuations du courant de fuite de surface, autre composante du bruit parallèle.
3. Bruit associé à une résistance série ou aux mauvais contacts électriques avec le détecteur, un contributeur au bruit série.

6.3 Modifications liées à la tension de polarisation du détecteur

Lorsque la tension de polarisation et le champ électrique sont faibles, la hauteur d'impulsion des rayonnements complètement arrêtés dans la couche d'appauvrissement continue d'augmenter avec la tension appliquée. Cette variation est causée par la collecte incomplète des porteurs de charge en raison du piégeage ou de la recombinaison le long de la trajectoire de la particule incidente. La

fraction qui échappe à la collection diminuera à mesure que le champ électrique augmente. Des pertes similaires dues à la recombinaison sont observées dans une chambre ionique remplie de gaz à de faibles valeurs du champ électrique. Une fois que le champ électrique est suffisamment élevé, la collecte des charges devient complète et la hauteur de l'impulsion ne change plus avec de nouvelles augmentations de la tension de polarisation du détecteur. Cette région de fonctionnement est appelée région de saturation et correspond à la région de saturation ionique dans une chambre ionique remplie de gaz.

6.4 Temps de montée des impulsions

Les détecteurs à diodes semi-conductrices sont généralement parmi les plus rapides de tous les détecteurs de rayonnement couramment utilisés. Dans des conditions normales, le temps de montée des impulsions observé est de l'ordre de 10 ns ou moins. La contribution du détecteur à ce temps de montée est composée du temps de transit de la charge et du temps de plasma.

6.5 Fenêtre d'entrée ou couche morte

Lorsqu'il s'agit de particules lourdement chargées ou d'autres rayonnements faiblement pénétrants, la perte d'énergie qui peut avoir lieu avant que la particule n'atteigne le volume actif du détecteur peut être importante. Étant donné que l'épaisseur de la couche morte comprend non seulement l'électrode métallique mais également une épaisseur indéterminée de silicium immédiatement sous l'électrode dans laquelle la collecte de charges est inefficace, la couche morte peut être fonction de la tension appliquée. Son épaisseur effective doit souvent être mesurée directement par l'utilisateur pour qu'une compensation précise puisse être effectuée.

6.6 Canalisation

Dans les matériaux cristallins, le taux de perte d'énergie d'une particule chargée peut dépendre de l'orientation de sa trajectoire par rapport aux axes cristallins. Les particules qui se déplacent parallèlement aux plans cristallins peuvent, en moyenne, présenter un taux de perte d'énergie qui est inférieure à celle des particules dirigées dans une direction arbitraire. Par conséquent, ces particules « canalisées » peuvent pénétrer beaucoup plus loin à travers le cristal. Les effets sont particulièrement significatifs pour les détecteurs minces totalement appauvris car la quantité d'énergie déposée dépend alors de l'orientation des plans cristallins par rapport à la direction des particules. Pour minimiser la tendance des particules incidentes à se canaliser, les détecteurs sont normalement fabriqués à partir de silicium découpé de manière à ce que l'orientation du cristal (111) soit perpendiculaire à la surface de la plaquette.

6.7 Dommages causés par les radiations

Le bon fonctionnement de tout détecteur à semi-conducteur dépend de la quasi-perfection du réseau cristallin pour éviter les défauts qui peuvent piéger les porteurs de charge et conduire à une collecte incomplète des charges. Toute utilisation intensive de ces détecteurs garantit toutefois que certains dommages du réseau se produiront en raison des effets perturbateurs du rayonnement mesuré lors de son passage à travers le cristal. Alors que l'énergie nécessaire à la création de paires électron-trou conduit à des processus entièrement réversibles qui ne laissent aucun dommage, les transferts d'énergie non ionisants vers les atomes du cristal conduisent à des changements irréversibles. Les effets ont tendance à être relativement mineurs pour les rayonnements légèrement ionisants (particules bêta ou rayons gamma), mais peuvent devenir très importants dans des conditions typiques d'utilisation pour des particules lourdes chargées. Par exemple, une exposition prolongée des détecteurs à barrière de surface en silicium à des ions lourds ou à des fragments de fission entraînera une augmentation mesurable du courant de fuite et une perte significative de la résolution énergétique du détecteur.

6.8 Calibrage énergétique

Lorsqu'ils sont appliqués à la mesure d'électrons rapides ou d'ions légers tels que des protons ou des particules alpha, les détecteurs de rayonnement à diode semi-conductrice répondent de manière très linéaire, et l'étalonnage énergétique obtenu pour un type de particule est très proche de celui obtenu en utilisant un type de rayonnement différent. Certaines observations montrent qu'il existe une petite différence dans la hauteur d'impulsion observée pour les protons et les particules alpha de même énergie, mais ces différences sont généralement de l'ordre de 1 % ou moins. Il n'est pas clair si ces différences reflètent une différence fondamentale dans l'énergie d'ionisation pour diverses particules ou si d'autres facteurs liés aux modes de perte d'énergie peuvent en être responsables. Pour toute application dans laquelle un étalonnage d'énergie absolue d'environ 1 % ou moins est requis, il est toujours préférable d'étalonner le détecteur en utilisant le même type de particule que celui impliqué dans la mesure elle-même.

6.9 Défaut de hauteur d'impulsion

La réponse des détecteurs à semi-conducteurs aux ions très lourds tels que les fragments de fission est moins simple. Il existe de nombreuses preuves que la hauteur d'impulsion observée pour les ions lourds est nettement inférieure à celle observée pour un ion léger de même énergie. Le défaut de hauteur d'impulsion est défini en unités d'énergie comme la différence entre l'énergie réelle de l'ion

lourd et son énergie apparente, telle que déterminée à partir d'un étalonnage énergétique du détecteur obtenu à l'aide de particules alpha.

7. Applications des détecteurs à semi-conducteurs à diodes de silicium

7.1 Spectroscopie générale des particules chargées

Depuis leur développement en tant que détecteurs pratiques au début des années 1960, les diodes au silicium sont devenues les détecteurs de choix pour la majorité des applications impliquant des particules lourdes chargées. La plupart des conclusions s'appliquent à d'autres spectroscopies de particules chargées impliquant des protons, des deutons ou d'autres ions lourds.

Les détecteurs à diode au silicium sont parfois utilisés pour la mesure des particules bêta et des électrons rapides, en particulier en tant que détecteurs minces à transmission totalement appauvrie.

7.2 Spectroscopie de particules alpha

Les diodes au silicium fonctionnant à température ambiante sont des détecteurs presque idéaux pour les particules alpha et autres ions légers. En raison de la grande disponibilité de sources mono-énergétiques pratiques de particules alpha, les performances des détecteurs à semi-conducteurs sont testées de manière classique en enregistrant les spectres de hauteur d'impulsion provenant de ces sources. Le plus courant d'entre eux est Am, et le spectre alpha correspondant est largement utilisé pour comparer la résolution énergétique des détecteurs à semi-conducteurs. Avec des particules alpha dans cette plage d'énergie (5,486 MeV), la contribution au bruit du préamplificateur et des autres composants électroniques peut être bien inférieure à la résolution énergétique inhérente du détecteur lui-même.

7.3 Spectroscopie d'ions lourds et de fragments de fission

La mesure de l'énergie des fragments de fission ou d'autres ions de grande masse implique plusieurs préoccupations particulières. La plupart proviennent de la haute densité de porteurs de charge créés le long du parcours de ces ions. La recombinaison des paires électron-trou est accentuée et le détecteur peut nécessiter une tension de polarisation supérieure à celle requise pour saturer le signal des particules alpha. Le défaut de hauteur d'impulsion évoqué précédemment est également accentué par cette forte densité de porteurs, compliquant les procédures d'étalonnage en énergie. En outre, l'exposition prolongée des détecteurs à des ions lourds ou à des fragments de fission crée une détérioration rapide des performances en raison des dommages causés par les radiations dans le détecteur.

Les fabricants de détecteurs proposent généralement des détecteurs au silicium spécialement adaptés à la spectroscopie des ions lourds. Ils sont conçus pour minimiser les problèmes de temps de montée lent et de défaut de hauteur d'impulsion causés par la perte d'énergie spécifique élevée le long de la trajectoire des particules. L'étape la plus efficace consiste à s'assurer que le champ électrique est aussi élevé que possible. Une approche consiste à utiliser de fines couches de silicium à faible résistivité pour préparer des détecteurs totalement appauvris qui nécessitent une grande valeur de polarisation inverse pour être complètement appauvris.

7.4 Mesures de perte d'énergie - Identification des particules

Précédemment, nous avons donné des exemples de spectroscopie de particules chargées qui impliquaient l'arrêt total de ces particules dans la région d'appauvrissement du détecteur à semi-conducteur. En négligeant le défaut de hauteur d'impulsion, le nombre de porteurs de charge créés est proportionnel à l'énergie totale du rayonnement incident. Il existe parfois des applications dans lesquelles la perte d'énergie spécifique dE/dx du rayonnement incident est intéressante, plutôt que son énergie totale. Pour ces applications, des détecteurs minces par rapport à la gamme de particules sont choisis. Le nombre de porteurs de charge créés au sein d'un détecteur de faible épaisseur Δt sera simplement $(dE/dx) \Delta t/E$. La particule traverse complètement le détecteur, conservant l'essentiel de son énergie initiale, et un signal proportionnel à dE/dx est observé. Dans de telles applications, l'appareil est souvent appelé détecteur ΔE .

7.5 Spectroscopie à rayons X avec diodes p-i-n au silicium

Même si le silicium a un numéro atomique relativement faible de 14, l'absorption photoélectrique des photons incidents est toujours prédominante dans la région des rayons X mous en dessous de 20 keV. Alors que les configurations à dérive de lithium et autres configurations de silicium sont devenues prédominantes pour la spectroscopie aux rayons X, une diode PIN en silicium plus simple peut constituer une alternative utile. Dans cette configuration, une région I à haute résistivité est dotée de contacts non injectant p et n sur chaque surface pour aider à réduire le courant de fuite en dessous de celui qui serait observé avec une simple diode. Une épaisseur typique de 300 micromètres est suffisante pour fournir une efficacité de détection utile jusqu'à 20 ou 30 keV. Le nombre relativement faible de paires électron-trou créées par un photon à rayon X de faible énergie nécessite que le niveau de bruit du détecteur soit réduit au minimum absolu.

7.6 Fonctionnement en mode photovoltaïque

Il est possible de mesurer le courant induit par un rayonnement à partir d'un détecteur à semi-conducteur, même en l'absence de tension appliquée. Ce mode de fonctionnement photovoltaïque est similaire à celui d'une cellule solaire ordinaire, sauf que les porteurs de charge sont créés par le rayonnement ionisant plutôt que par la lumière incidente. La jonction de semi-conducteurs de type n et p crée un potentiel de contact et une région d'appauvrissement. Les électrons et les trous formés dans la région d'appauvrissement dériveront sous l'influence du champ électrique correspondant et le courant sera mesuré dans un circuit externe connecté aux bornes de la jonction. Ce mode photovoltaïque est peu utilisé car le potentiel de contact est faible, inférieur à 1 V dans le silicium. L'épaisseur de la région d'appauvrissement est également faible, généralement inférieure à la gamme de particules chargées d'intérêt. Les détecteurs à jonction fonctionnent donc normalement avec une tension de polarisation inverse élevée pour améliorer leur efficacité de collecte de charges et augmenter l'épaisseur de leur volume actif. Cependant, la petite région d'appauvrissement présente sans tension appliquée peut être suffisante pour produire un courant utile lorsque l'intensité du rayonnement incident est suffisamment élevée.

7.7 Diodes de silicium comme moniteurs de personnel

Les diodes au silicium, généralement de configuration PIN, sont devenues disponibles dans le commerce sous forme de capteurs de rayonnement compacts à porter comme moniteurs de personnel. Comparés aux moniteurs passifs tels que les dosimètres thermo-luminescents, ils offrent l'avantage d'une lecture en temps réel. D'autres détecteurs actifs tels que les tubes GM ont été utilisés pour une surveillance similaire, mais les diodes au silicium offrent un détecteur de rayonnement plus compact et plus robuste. Étant donné que les diodes sur lesquelles elles sont basées peuvent avoir de petites dimensions, des conceptions ont émergé avec des formats de poche pratiques. Les instruments disponibles dans le commerce sont souvent commercialisés sous le nom de « dosimètre personnel électronique » ou EPD. Lorsqu'elle est exposée aux rayons X ou aux rayons gamma, la diode au silicium produira des impulsions à partir d'électrons secondaires créés soit dans le volume actif de la diode, soit dans les matériaux environnants à partir desquels les électrons secondaires atteignent le volume actif du détecteur.

Lorsqu'elle est exposée aux rayons X ou aux rayons gamma, la diode de silicium produira des impulsions à partir d'électrons secondaires qui sont créés soit dans le volume actif de la diode, soit dans des matériaux environnants à partir desquels les électrons secondaires atteignent le volume actif du détecteur. Idéalement, un nombre enregistré à partir de la diode devrait représenter une unité d'exposition, indépendamment de l'énergie de la radiographie incidente ou du rayon gamma.

Annexe III

B - Détecteurs au germanium

1. Considérations générales

Les détecteurs à barrière de surface et les diodes à implantation ionique trouvent une application répandue pour la détection de particules alpha et d'autres rayonnements à courte portée, mais ne sont pas facilement adaptables à des applications impliquant des rayonnements plus pénétrants. Leur majeure limitation est la profondeur d'épuisement maximale ou le volume actif pouvant être créé. En utilisant du silicium ou du germanium de pureté semi-conductrice normale, des profondeurs d'appauvrissement supérieures à 2 ou 3 mm sont difficiles à atteindre malgré l'application de tensions de polarisation proches du niveau de claquage. Des épaisseurs bien plus importantes sont nécessaires pour les détecteurs destinés à la spectroscopie gamma, qui font l'objet de cette annexe. De l'équation (28) l'épaisseur de la région d'appauvrissement est donnée par :

$$d = \left(\frac{2\epsilon V}{eN} \right)^{1/2} \quad (32)$$

À une tension appliquée donnée, des profondeurs d'épuisement plus importantes ne peuvent être obtenues qu'en diminuant la valeur de N grâce à des réductions supplémentaires de la concentration nette d'impuretés.

Deux approches générales peuvent être adoptées pour atteindre cet objectif. La première consiste à rechercher des techniques de raffinage plus avancées capables de réduire la concentration en impuretés à environ 10^{10} atomes / cm^3 . Plusieurs techniques ont été développées pour atteindre cet objectif en germanium, mais pas en silicium. Les détecteurs fabriqués à partir du germanium ultra-pur sont généralement appelés détecteurs au germanium intrinsèque ou au germanium de haute pureté (HPGe).

La deuxième approche pour réduire la concentration nette d'impuretés consiste à créer un matériau compensé dans lequel les impuretés résiduelles sont équilibrées par une concentration égale d'atomes dopants de type opposé. Le processus de dérive des ions lithium a été appliqué aux cristaux de silicium et de germanium pour compenser le matériau après la croissance du cristal. Le matériau compensé résultant possède de nombreuses propriétés d'un matériau intrinsèque ou pur.

Les détecteurs au germanium produits par le procédé de dérive du lithium reçoivent la désignation Ge(Li). Ils ont servi de type courant de détecteur au germanium à grand volume pendant deux décennies. La disponibilité généralisée de germanium de haute pureté a fourni une alternative à la dérive du

lithium, et les fabricants ont désormais abandonné la production de détecteurs Ge(Li) au profit du type HPGe. L'une des principales raisons de cette évolution, est la commodité opérationnelle bien plus grande offerte par les détecteurs HPGe. Alors que les détecteurs Ge(Li) doivent être maintenus en permanence à basse température, les détecteurs HPGe peuvent être laissés se réchauffer à température ambiante entre deux utilisations.

En général, les caractéristiques de performances importantes telles que l'efficacité de détection et la résolution énergétique sont essentiellement identiques pour les détecteurs Ge(Li) et HPGe de même taille.

2. Configurations de détecteurs au Germanium

2.1 Fabrication de détecteurs au germanium de haute pureté (HPGe)

Dans les techniques de production de germanium ultra pur avec des niveaux d'impuretés aussi faibles que 10^{10} atomes/cm³ le matériau de départ est du germanium en vrac destiné à être utilisé dans l'industrie des semi-conducteurs. Ce matériau, déjà d'une grande pureté, est ensuite traité selon la technique du raffinage de zone. Les niveaux d'impuretés sont progressivement réduits de chauffer localement le matériau et faire passer lentement une zone fondue d'une extrémité à l'autre de l'échantillon. Puisque les impuretés ont tendance à être plus solubles dans le germanium fondu que dans le solide, les impuretés sont préférentiellement transférées vers la zone fondue et sont balayées de l'échantillon. Après de nombreuses répétitions de ce processus, des niveaux d'impuretés aussi bas que 10^9 atomes/cm³ seront atteints. Si les impuretés restantes de faible niveau sont des accepteurs (tels que l'aluminium), les propriétés électriques du cristal semi-conducteur issu du matériau sont légèrement de type p. Alternativement, si les impuretés du donneur demeurent, un type n de haute pureté est le résultat.

Les détecteurs au germanium pratiques sont toujours fabriqués à partir d'un matériau dans lequel la concentration nette de dopant est très faible et la conductivité est donc faible. Dans la nomenclature plus largement utilisée, cette condition est généralement identifiée comme étant un « matériau à haute résistivité ».

2.2 Configuration planaire

Une configuration représentative d'un détecteur HPGe planaire fabriqué à partir de germanium de type p de haute pureté est présentée sur la figure 12. Dans une configuration planaire, les contacts électriques sont prévus sur les deux surfaces planes d'un disque en germanium. Le contact n⁺ peut être formé soit par évaporation et diffusion du lithium sur une surface de la plaquette, soit par implantation directe d'atomes donneurs à l'aide d'un accélérateur. La région d'appauvrissement du

détecteur est formée par polarisation inverse de cette jonction. Le contact sur la face opposée du cristal doit être un contact non injectant pour un porteur majoritaire. Il peut s'agir d'un contact p+ produit par implantation ionique d'atomes accepteurs, ou d'une barrière superficielle métal-semi-conducteur qui agit comme son équivalent électrique. Les techniques d'implantation ionique présentent l'avantage que les couches de contact peuvent être très fines pour servir de fenêtres d'entrée aux rayonnements faiblement pénétrants tels que les rayons X de faible énergie. La formation de contacts de type p est simple par implantation de bore. La production de contacts de type n, est moins courante puisque les dommages par radiation provoqués par le processus d'implantation produisent des sites accepteurs dans le germanium. Les contacts de type n peuvent être établis par implantation de phosphore, mais une étape de recuit difficile est nécessaire pour garantir que ces atomes donneurs deviennent électriquement actifs. Par conséquent, la majorité des contacts de type n sont toujours produits par diffusion de lithium, ce qui entraîne une couche morte superficielle un peu plus épaisse.

Les détecteurs au germanium de haute pureté fonctionnent généralement comme des détecteurs entièrement épuisés. La polarisation inverse nécessite qu'une tension positive soit appliquée au contact n+ par rapport à la surface p+. La région d'appauvrissement commence effectivement au bord n+ de la région centrale et s'étend plus profondément à mesure que la tension augmente jusqu'à ce que le détecteur soit complètement épuisé. Des augmentations supplémentaires de la tension servent alors à augmenter le champ électrique partout dans le détecteur d'une quantité uniforme.

Dans le germanium, à la température normale de fonctionnement de 77 K, les vitesses saturées des électrons sont atteintes avec un champ minimum d'environ 10^5 V/m, mais des intensités de champ trois à cinq fois plus grandes sont nécessaires pour saturer complètement la vitesse des trous. La tension maximale (généralement 3 à 5 kV) valeurs auxquelles les électrons, mais pas les trous, atteindront une vitesse de dérive saturée.

Les détecteurs au germanium peuvent également être fabriqués à partir d'un matériau de type n de haute pureté plutôt qu'avec un matériau de type p comme décrit ci-dessus. Dans ce cas, des contacts n+ et p+ sont à nouveau prévus sur chaque surface de la tranche. La polarisation inverse nécessite toujours l'application d'une tension positive au contact n+ par rapport à la surface p+. Toutefois, les rôles des deux interlocuteurs sont inversés. La couche p+ sert maintenant de contact redresseur et la région d'appauvrissement s'étend à partir de ce contact à mesure que la tension augmente. Le contact n+ fonctionne désormais comme un contact bloquant ou non injectant dans lequel la population de trous est très faible.

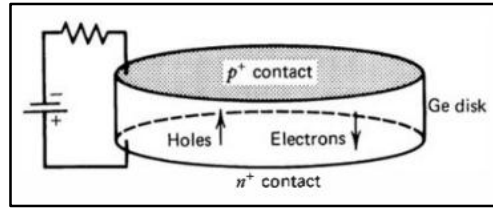


Fig. 12 Configuration d'un détecteur planaire HPGe.

Un troisième type de configuration de diode représente le cas des détecteurs à dérivate de lithium. Dans ce cas, le processus de dérivate du lithium est utilisé pour convertir la majeure partie de la tranche en matériau compensé ou intrinsèque. Ici, les concentrations de donneurs et d'accepteurs s'équilibrent exactement, et le matériel n'est ni de type n ni de type p. Là encore, des contacts n^+ et p^+ sont prévus sur chaque face de la tranche. Dans ce cas, il n'y a pas de distinction entre le contact redresseur et le contact de blocage. Si la dérivate du lithium est encore utilisée pour compenser le silicium, elle n'est plus réalisée dans le germanium depuis l'introduction du matériau ultra-pur utilisé dans les détecteurs HPGe.

Lorsqu'il est complètement épuisé et fonctionne avec une surtension importante, le champ électrique des détecteurs planaires est presque uniforme d'un contact à l'autre. Les trous ou électrons dérivent donc sous l'influence d'un champ quasi constant dans tout le volume actif du détecteur. Dans ces circonstances, les configurations décrites ci-dessus présentent des propriétés très similaires à celles des détecteurs à géométrie planaire. En revanche, le champ électrique est assez non uniforme dans les détecteurs à géométrie coaxiale décrits ci-dessous et le choix de la configuration a donc une influence plus significative sur les détails du mouvement des porteurs de charge.

Les détecteurs planaires au germanium et les détecteurs au silicium à dérivate de lithium sont généralement beaucoup plus épais (plusieurs millimètres) que les diodes au silicium standard. Par conséquent, une tension de fonctionnement sensiblement plus élevée d'au moins 1 kV ou plus est nécessaire pour développer l'intensité du champ électrique nécessaire pour collecter efficacement les électrons et les trous. Dans la configuration planaire simple illustrée à la figure (12), la pleine tension appliquée apparaît sur les surfaces latérales du disque et peut conduire à un courant de fuite important. Les fluctuations de ce courant sont une source de bruit et peuvent conduire à une détérioration de la résolution énergétique du détecteur.

Une méthode permettant de supprimer considérablement les effets de ce courant de fuite est illustrée à la figure (13a). Il utilise une électrode à anneau de garde étroite sur la surface de lecture autour de la périphérie de l'électrode de lecture, et les deux sont séparées par une petite espace. Le courant de fuite circulant à travers les surfaces latérales est simplement conduit vers la terre et n'influence plus directement le signal mesuré. Étant donné que l'entrée du préamplificateur est

également une masse virtuelle, toute fuite restante à travers l'espace qui influence la lecture est considérablement réduite. L'efficacité du détecteur est légèrement pénalisée puisque la surface de l'électrode de lecture est légèrement plus petite qu'elle ne le serait si elle couvrait toute la surface.

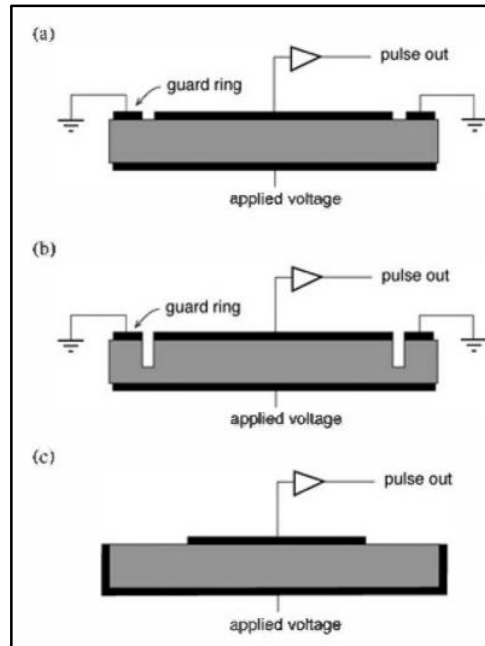


Fig. 13 Trois variantes sur la configuration des détecteurs planaires. Les électrodes sont représentées avec une épaisseur exagérée sous forme de lignes noires continues.

(a) Une électrode à anneau de garde étroit protège l'électrode de lecture des effets du courant de fuite conduit à travers la surface latérale cylindrique du disque semi-conducteur.

(b) Une rainure est prévue dans le volume semi-conducteur pour augmenter la surface séparant l'anneau de garde et l'électrode de lecture.

(c) Une électrode « enveloppante » est étendue à partir d'une surface plane pour couvrir également la surface latérale du disque, et une électrode de lecture de petit diamètre est fabriquée sur la surface plane opposée. Le détecteur est normalement orienté de manière à ce que les rayons gamma ou X soient incidents sur la surface indiquée en bas dans ce schéma.

D'un autre côté, le volume actif exclut désormais les régions du semi-conducteur proches des bords du disque où un certain piégeage de charges peut se produire, conduisant à des impulsions d'amplitude inférieure à celle appropriée. La figure (13b) montre également une configuration rainurée qui utilise une tranchée fabriquée dans le volume semi-conducteur pour augmenter la surface effective de l'espace sans réduire davantage le volume actif du détecteur.

Une autre géométrie parfois utilisée pour les détecteurs planaires au germanium est illustrée sur la figure (13c). Ici, l'électrode représentée sur la surface inférieure est « enroulée autour » de la surface cylindrique latérale du disque, et l'électrode supérieure est considérablement plus petite que la surface plane complète. Ce choix offre un large écart entre les deux électrodes et présente l'avantage supplémentaire de réduire la capacité du détecteur en dessous de la valeur qu'elle aurait si l'électrode

de lecture couvrait toute la surface. Étant donné que le bruit électronique évolue généralement avec la capacité du détecteur, il existe un avantage potentiel en termes de rapport signal/bruit lorsqu'il s'agit de petites impulsions provenant de rayons gamma ou de rayons X de faible énergie. Dans cette géométrie, les lignes de champ électrique ne sont plus parallèles comme dans un simple détecteur plan, de sorte que les analyses de l'intensité du champ électrique et des formes d'impulsion qui suivent pour une géométrie plane simple ne sont plus strictement applicables. Étant donné que l'application courante de ces détecteurs plans est la spectroscopie des rayons gamma et X de faible énergie, l'électrode enveloppante est maintenue aussi fine que possible pour minimiser l'atténuation, et le disque est normalement orienté de manière à ce que cette fine électrode fasse face à la fenêtre d'entrée du boîtier du détecteur.

2.3 Configuration coaxiale

Pour les détecteurs plans décrits ci-dessus, le volume total de germanium est généralement limité à des dizaines de cm^3 par le diamètre du cristal développé et l'épaisseur de la tranche découpée. Pour produire des détecteurs avec un volume plus grand, comme cela est préféré pour la spectroscopie générale des rayons gamma, une géométrie différente est choisie qui peut incorporer une plus grande partie du volume cristallin dans un seul détecteur fini. Une géométrie cylindrique ou coaxiale est maintenant choisie comme indiqué sur la figure (14). Dans ce cas, une électrode est fabriquée au niveau de la surface cylindrique externe d'un long cristal cylindrique de germanium. Le deuxième contact cylindrique est obtenu en retirant le noyau du cristal et en plaçant un contact sur la surface cylindrique intérieure. Étant donné que le cristal peut être allongé dans la direction axiale, des volumes actifs beaucoup plus importants peuvent être produits. Un avantage supplémentaire de la géométrie coaxiale est que, en utilisant un petit diamètre intérieur, des détecteurs de grand volume peuvent être fabriqués avec une capacité inférieure à celle qui serait possible en utilisant une géométrie plane.

Une configuration coaxiale à extrémités fermées est une configuration dans laquelle seule une partie du noyau central est retirée et l'électrode externe s'étend sur une extrémité plate du cristal cylindrique (voir Fig. 14).

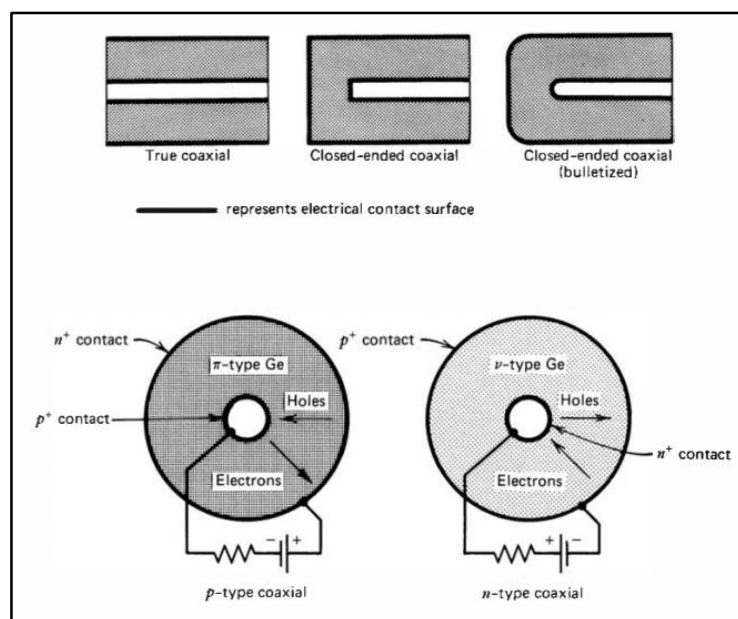


Fig. 14 En haut sont représentées les trois formes courantes de détecteurs coaxiaux à grand volume. Chacune représente une vue en coupe coaxiale passant par l'axe d'un cristal cylindrique. L'électrode externe s'étend sur la surface plane avant (gauche) dans les deux cas à extrémités fermées. Les sections transversales perpendiculaires à l'axe cylindrique du cristal sont les contacts indiqués en bas.

Le matériau HPGe peut être de type p ou n de haute pureté. Les configurations d'électrodes correspondantes sont indiquées pour chaque type.

La plupart des fabricants commerciaux de détecteurs HPGe choisissent la configuration fermée plutôt qu'une véritable géométrie coaxiale pour éviter les complications liées à la gestion des courants de fuite au niveau de la surface avant. En outre, la configuration à extrémités fermées fournit une surface avant plane qui peut servir de fenêtre d'entrée pour les rayonnements faiblement pénétrants si elle est fabriquée avec un mince contact électrique. Dans la configuration fermée, les lignes de champ électrique ne sont plus complètement radiales comme elles le sont dans le véritable boîtier coaxial, et il y a une certaine tendance à produire des régions d'intensité de champ réduite près des coins du cristal où les vitesses des porteurs peuvent être inférieures à normale. L'extension du trou central pour atteindre près de la surface avant permet de maintenir les lignes de champ aussi radiales que possible, et la plupart des fabricants arrondissent les coins avant du cristal et du trou (appelé bulletizing) pour aider à éliminer les régions à faible champ.

Les détecteurs coaxiaux HPGe sont également disponibles dans des configurations de puits dans lesquelles le boîtier est façonné pour permettre un accès externe au trou central. De petites sources de radio-isotopes peuvent ensuite être placées dans ce puits pour des mesures dans lesquelles la source est presque entourée de germanium et l'efficacité de détection peut être inhabituellement élevée.

En géométrie coaxiale, le contact redresseur qui forme la jonction semi-conductrice peut en principe être placé soit à la surface interne, soit à la surface externe du cristal. Les conditions de champ électrique qui en résultent sont très différentes. Si le contact redresseur se trouve sur la surface extérieure, la couche d'appauvrissement se développe vers l'intérieur à mesure que la tension augmente jusqu'à atteindre la surface du trou intérieur à la tension d'appauvrissement. Si la surface interne comporte un contact redresseur, la région d'appauvrissement s'étend vers l'extérieur et une tension beaucoup plus élevée est nécessaire pour épuiser complètement le volume du détecteur. Le premier choix maintient également une valeur de champ électrique plus élevée dans les régions extérieures du cristal cylindrique (où se trouve la majeure partie du volume) lorsque la tension est élevée au-dessus de l'épuisement complet.

2.4 Champ électrique et capacité

Le champ électrique dans les détecteurs au germanium détermine la vitesse de dérive des porteurs de charge. Sa configuration est donc importante pour les considérations relatives à la forme des impulsions, au comportement temporel et à l'exhaustivité du processus de collecte de charges. La variation spatiale de l'intensité du champ à travers le volume actif du détecteur est nettement différente en géométrie planaire et coaxiale, et les deux cas sont considérés séparément dans la discussion suivante. Dans chaque géométrie, on résout l'équation de Poisson :

$$\boxed{\nabla^2 \varphi = -\frac{\rho}{\epsilon}} \quad (33)$$

Pour trouver le potentiel électrique φ en présence de la densité de charge ρ (ϵ est la constante diélectrique). La densité de charge dépend du matériau à partir duquel le détecteur est fabriqué.

Pour le germanium de type p, $\rho = -eN_A$, où e est la charge électronique et N_A est la densité des impuretés accepteurs. Cette charge négative représente les sites accepteurs remplis du côté informatique de la jonction. De même, les sites donneurs ionisés dans le germanium de type V représentent une densité de charge positive donnée par $\rho = eN_D$. Dans le germanium dérivé du lithium, des impuretés de chaque type sont exactement compensés, et $\rho = 0$. Dans les dérivations qui suivent, nous énonçons explicitement les résultats pour le premier cas de germanium de type n. Les résultats correspondants pour les détecteurs de type p et à dérive de lithium peuvent être obtenus en substituant la valeur appropriée à ρ .

2.4.1 Géométrie planaire

D'après les résultats dérivés de l'annexe III-A pour les diodes planaires (Eq. 10), la profondeur d'appauvrissement du détecteur est :

$$d = \left(\frac{2\epsilon V}{\rho} \right)^{1/2} \quad (34)$$

L'épuisement (déplétion) complet nécessite une tension appliquée minimale V_d (la tension de déplétion) à laquelle la profondeur d'épuisement s'étend entièrement sur l'épaisseur de la tranche T :

$$V_d = \frac{\rho T^2}{2\epsilon} \quad (35)$$

Dans la géométrie des dalles unidimensionnelles, l'équation de Poisson (Eq. 33) devient :

$$\frac{d^2\phi}{dx^2} = -\frac{\rho}{\epsilon} \quad (36)$$

Pour une tension appliquée inférieure à celle requise pour un épuisement complet, le champ électrique $\xi = -\text{grad } \phi = -d\phi / dx$ est obtenu par solution de l'équation (36) avec la condition aux limites :

$$\phi(d) - \phi(0) = V \quad (37)$$

Le résultat est :

$$\left[-\mathcal{E}(x) = \frac{V}{d} + \frac{\rho}{\epsilon} \left(\frac{d}{2} - x \right) \right] \left[|\mathcal{E}(x)| = \frac{V}{d} + \frac{eN_A}{\epsilon} \left(x - \frac{d}{2} \right) \right] \quad (38)$$

Où x est la distance du contact p^+ (voir Fig. 15). Pour $V < V_d$, la partie de cette solution correspondant à la région non appauvrie du détecteur n'est pas applicable, et le champ est nul. L'équation (38) est également valable pour $V > V_d$, car l'effet de la surtension $(V - V_d)$ est d'augmenter le champ d'une quantité constante $(V - V_d) / T$ partout dans le détecteur.

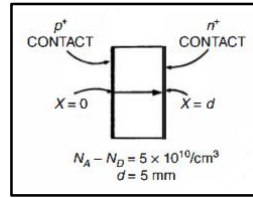


Fig. 15 Volume actif d'un détecteur planaire HPGe

Comme dans les détecteurs à barrière de surface ou à jonction, la capacité varie avec la tension de polarisation appliquée jusqu'à la valeur à laquelle le détecteur est complètement épuisé. En géométrie planaire, la capacité du détecteur par unité de surface avant le point d'épuisement complet est donnée par :

$$C = \left(\frac{\epsilon \rho}{2V} \right)^{1/2} \quad (39)$$

Pour $V > V_d$, la capacité du détecteur est une constante obtenue en réglant $V = V_d$ dans l'équation (39). L'indépendance de la capacité du détecteur par rapport à la polarisation appliquée est souvent considérée comme une indication d'un épuisement complet du détecteur.

2.4.2 Géométrie coaxiale

L'équation de Poisson (Eq. 33) en coordonnées cylindriques devient :

$$\frac{d^2 \phi}{dr^2} + \frac{1}{r} \frac{d\phi}{dr} = -\frac{\rho}{\epsilon} \quad (40)$$

Nous traitons le cas d'un véritable détecteur coaxial avec des rayons intérieur et extérieur de r_1 et r_2 , (voir Fig. 16). Une condition aux limites est que la différence de potentiel entre ces rayons soit donnée par la tension appliquée V , ou $\phi(r_2) - \phi(r_1) = V$. La résolution de l'équation. (40) pour $\xi = -d\phi / dr$, nous trouvons que la configuration du champ électrique résultante est :

$$\left[-\mathcal{E}(r) = -\frac{\rho}{2\epsilon} r + \frac{V + (\rho/4\epsilon)(r_2^2 - r_1^2)}{r \ln(r_2/r_1)} \right] \left[|\mathcal{E}(r)| = \frac{eN_A}{2\epsilon} r + \frac{V - (eN_A/4\epsilon)(r_2^2 - r_1^2)}{r \ln(r_2/r_1)} \right] \quad (41)$$

à condition que N_A (la concentration de l'accepteur) soit constante sur le volume du détecteur.

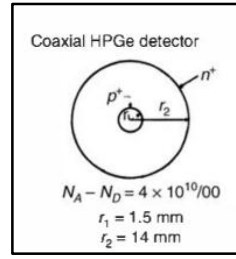


Fig. 16 Volume actif d'un détecteur coaxial HPGe

La valeur absolue du profil de champ électrique reste la même pour les configurations coaxiales de type p et de type n illustrées sur la figure (14), à condition que la même tension et les mêmes concentrations de dopants soient impliquées. Seule la polarité du champ change, inversant les directions de dérive des trous et des électrons. Notez que les deux cas maintiennent le contact redresseur au niveau de la surface extérieure.

La tension V_d nécessaire pour épuiser complètement le détecteur peut être trouvée en posant $\xi(r_1) = 0$ dans l'équation (41) ci-dessus. Le résultat est :

$$V_d = \frac{\rho}{2\epsilon} \left[r_1^2 \ln\left(\frac{r_2}{r_1}\right) - \frac{1}{2}(r_2^2 - r_1^2) \right] \quad (42)$$

La tension d'appauvrissement diminue linéairement avec ρ (ou niveau de dopant) et est théoriquement nulle pour le germanium parfaitement compensé.

La capacité par unité de longueur d'un véritable détecteur coaxial entièrement appauvri est :

$$C = \frac{2\pi\epsilon}{\ln(r_2/r_1)} \quad (43)$$

Parce qu'il y a généralement un avantage à minimiser la capacité du détecteur, le rayon du noyau central r_1 est normalement maintenu à un minimum compatible avec la fabrication d'un contact électrique approprié.

2.5 Couche morte en surface

En première approximation, le volume actif d'un détecteur au germanium est simplement la région située entre les contacts n^+ et p^+ . Cependant, ces contacts peuvent avoir une épaisseur appréciable et peuvent représenter une couche morte à la surface du cristal à travers laquelle doit passer le rayonnement incident. Par exemple, certaines méthodes traditionnelles de fabrication de contacts, telles que l'évaporation et la diffusion du lithium dans la surface pour former une couche n^+ , peuvent produire des épaisseurs de plusieurs centaines de micromètres. Pour les rayons gamma d'une énergie d'environ 200 keV ou plus, l'atténuation dans de telles couches est généralement

négligeable et l'efficacité de détection des rayons gamma n'est pas sensiblement affectée par la présence de la couche morte. Cependant, si des rayons gamma ou des rayons X de plus faible énergie doivent être mesurés, la présence de couches de cette épaisseur doit être évitée pour éviter toute atténuation.

Pour de telles applications, la technique d'implantation ionique est devenue plus populaire pour la formation des contacts. Par exemple, le contact p^+ peut être produit en implantant des ions bore accélérés jusqu'à environ 20 keV dans la surface du cristal. De tels contacts n'ont que quelques dixièmes de micromètres d'épaisseur et peuvent servir de fenêtres d'entrée appropriées pour les rayons X mous. Un détecteur au germanium de grand volume fabriqué avec un mince contact de surface implanté d'ions sera donc utile pour la mesure des photons sur une large gamme d'énergie, depuis les rayons X de quelques keV jusqu'aux rayons gamma de plusieurs MeV d'énergie.

La couche morte de surface des détecteurs au germanium peut varier lentement au fil du temps en raison de la formation de ce que l'on appelle les canaux de surface dans lesquels le champ électrique et l'efficacité de la collecte des charges sont réduits. Il est donc conseillé de répéter périodiquement les mesures d'efficacité, en particulier aux énergies de rayons gamma plus faibles, où ces effets seront les plus significatifs.

3. Caractéristiques opérationnelles du détecteur de germanium

3.1 Cryostat et Dewar du détecteur

En raison de la petite bande interdite (0,7 eV), le fonctionnement à température ambiante de tout type de détecteur au germanium est impossible en raison de l'important courant de fuite induit thermiquement qui en résulterait. Au lieu de cela, les détecteurs au germanium doivent être refroidis pour réduire le courant de fuite au point que le bruit associé ne gâche pas leur excellente résolution énergétique. Normalement, la température est réduite à 77 K grâce à l'utilisation d'un Dewar isolé dans lequel un réservoir d'azote liquide est maintenu en contact thermique avec le détecteur.

3.2 Résolution énergétique

La caractéristique dominante des détecteurs au germanium est leur excellente résolution énergétique lorsqu'ils sont appliqués à la spectroscopie gamma. La comparaison des spectres de hauteur d'impulsion pour un scintillateur NaI(Tl) et un détecteur au germanium pour des spectres de rayons gamma incidents identiques montre la nette supériorité du système germanium en résolution énergétique permettant la séparation de nombreuses énergies de rayons gamma rapprochées, qui restent non résolues dans le spectre NaI(Tl).

Par conséquent, pratiquement toutes les spectroscopies de rayons gamma impliquant des spectres d'énergie complexes sont désormais réalisées avec des détecteurs au germanium.

3.3 Forme d'impulsion et propriétés de synchronisation

3.3.1 Processus de collecte de charges

La forme détaillée de l'impulsion du signal développée par un détecteur au germanium est importante à plusieurs égards. Dans la spectroscopie standard de hauteur d'impulsion, il est nécessaire que les temps de mise en forme de l'électronique de traitement des impulsions soient sensiblement supérieurs au temps de montée le plus long susceptible d'être rencontré par le détecteur si l'on veut éviter une perte de résolution due à un déficit balistique. Dans les applications où les informations de synchronisation doivent être obtenues à partir de l'impulsion, le temps de montée et la forme détaillée du front montant de l'impulsion deviennent importants lorsque l'on considère diverses méthodes de détection de temps. La résolution temporelle ultime pouvant être obtenue à partir des détecteurs au germanium dépend de manière critique à la fois du temps de montée moyen global et de la variation significative de la forme de l'impulsion d'un événement à l'autre.

En supposant que le circuit équivalent de l'électronique de mesure présente une constante de temps élevée par rapport au temps de montée le plus important produit par le détecteur, le front montant de l'impulsion du signal est presque entièrement déterminé par les détails du processus de collecte de charges au sein du détecteur. Etant donné qu'il est toujours avantageux d'avoir le temps de montée le plus court possible, on recherche normalement des conditions dans lesquelles la collecte des charges se produit dans le temps le plus court possible.

Deux facteurs limitent la résolution temporelle pouvant être obtenue avec les détecteurs au germanium. La première est que le processus de perception des charges est intrinsèquement lent. Même à la vitesse de dérive saturée, le temps nécessaire à un porteur de charge pour parcourir une distance de 1 cm (qui peut être de l'ordre de l'épaisseur du détecteur) est d'environ 100 ns. Les temps typiques de montée des impulsions peuvent donc atteindre plusieurs centaines de nanosecondes. Ces temps sont beaucoup plus longs que ceux des détecteurs rapides (tels que les scintillateurs organiques) et les performances temporelles ne pourront donc jamais être aussi bonnes. Cependant, un deuxième facteur aggrave encore les propriétés temporelles. Comme cela sera illustré ci-dessous, la forme détaillée de la montée d'impulsion des détecteurs au germanium peut changer considérablement d'un événement à l'autre, en fonction de la position à laquelle les paires électron-trou sont créées dans le volume actif. Lorsque le volume du détecteur est uniformément irradié, ces positions sont réparties de manière plus ou moins aléatoire, et donc les impulsions de sortie présentent une grande variation dans la forme de leur front montant.

3.3.2 Modèles pour la forme d'impulsion

Tout comme dans d'autres détecteurs dans lesquels les porteurs de signaux doivent être collectés sur des distances appréciables, la forme du front avant des impulsions des détecteurs au germanium dépend de la position à laquelle les porteurs de charge se forment dans le volume actif. Le cas le plus simple est celui d'une particule à très courte portée qui, en première approximation, crée toutes les paires électron-trou en un seul endroit du détecteur. (Pour les électrons secondaires créés par des rayons gamma d'énergies typiques, les longueurs de piste ne dépassent pas un millimètre ou deux et ne représentent généralement qu'une petite fraction des dimensions du détecteur.) Si cet emplacement se trouve à l'intérieur du volume actif, il y aura des temps de collecte uniques et séparés pour les trous et les électrons car chaque espèce doit parcourir une distance fixe avant d'être collectée. Si le point d'interaction est proche de l'un ou l'autre bord du volume actif, l'augmentation de l'impulsion observée sera alors principalement due au mouvement d'un seul type de porteur de charge. Pour les rayonnements chargés dont la portée n'est pas petite par rapport au volume actif, une distribution des temps de collecte résultera de la distribution spatiale correspondante des points de formation des trous et des électrons. Si l'orientation de la trajectoire des particules peut changer de manière significative d'un événement à l'autre, une variation supplémentaire du temps de montée de l'impulsion sera introduite.

La forme du front montant de l'impulsion de sortie est soumise à peu près à la même analyse que celle donnée pour les chambres à ions (détecteurs à gaz). Une différence majeure par rapport aux compteurs remplis de gaz est que les mobilités des porteurs de charge positifs et négatifs (trous et électrons) sont quelque peu similaires, alors que dans un gaz, la mobilité des ions positifs est bien inférieure à celle des électrons libres. Si un nombre suffisant d'hypothèses simplificatrices sont formulées, des expressions analytiques peuvent être dérivées pour décrire la forme de l'impulsion attendue dans des configurations simples de détecteurs au germanium. Les hypothèses courantes sont les suivantes :

- Tous les porteurs de charge sont créés à une position fixe dans le volume actif du détecteur.
- Le piégeage et le dépiégeage des porteurs de charge sont ignorés.
- Tous les porteurs de charge sont supposés être entièrement générés dans le volume actif du détecteur où le champ électrique a sa pleine valeur attendue.
- Le champ électrique dans le volume actif du détecteur est suffisamment élevé pour provoquer une saturation de la vitesse de dérive des électrons et des trous.

Avec cet ensemble d'hypothèses simplificatrices, des expressions analytiques peuvent être obtenues pour la forme d'impulsion attendue pour deux géométries relativement simples : le détecteur

plan dans lequel le champ électrique est approximativement uniforme, et la configuration coaxiale réelle dans laquelle le champ électrique varie inversement avec la distance radiale de l'axe du détecteur.

a. Géométrie planaire

Le développement de l'impulsion du signal peut ainsi être représenté par la croissance d'une charge induite $Q(t)$ dépendant du temps. Il existe deux composantes distinctes de la charge induite, l'une correspondant au mouvement des électrons et l'autre au mouvement des trous.

$$Q(t) = \frac{q_0}{d} (\text{electron drift distance} + \text{hole drift distance}) \quad (44)$$

Cette charge induite commence à zéro lorsque les électrons et les trous sont formés pour la première fois par la particule ionisante et atteint son maximum de q_0 lorsque les deux espèces ont été collectées. Comme dans le cas de la chambre ionique, $q_0 = n_0 e$, où n_0 est maintenant le nombre de paires électron-trou et e est la charge électronique.

Nous supposons maintenant que ces charges se forment à une distance x du contact n^+ du détecteur plan. Les définitions suivantes sont utilisées :

t_e = electron collection time
= x/v_e , where v_e is the saturation velocity

t_h = hole collection time
= $\frac{d-x}{v_h}$, where v_h is the saturation hole velocity

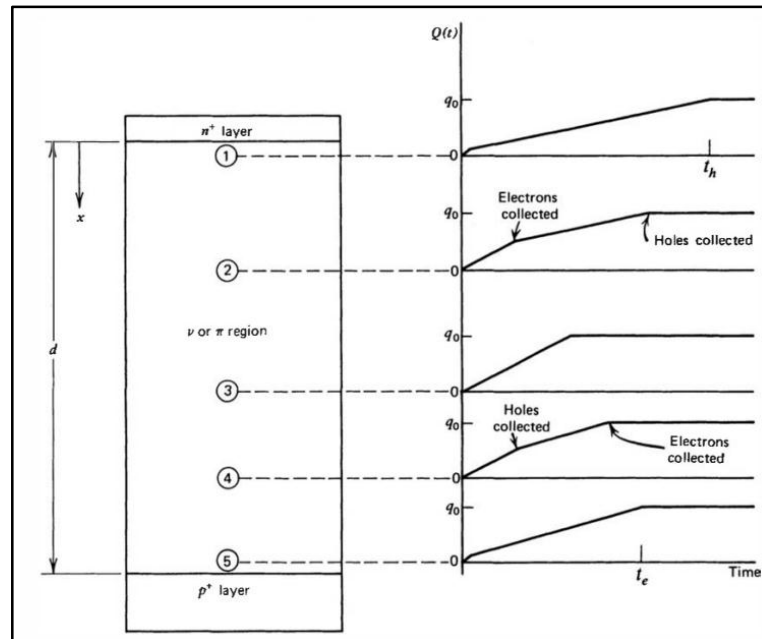


Fig. 17 Forme du front montant de l'impulsion de sortie $Q(t)$ pour différents points d'interaction au sein du détecteur

L'équation (44) peut être divisée en quatre domaines temporels possibles comme suit. Pendant que les trous et les électrons dérivent ($t < t_h$ et $t < t_e$) :

$$Q(t) = q_0 \left(\frac{v_e}{d} t + \frac{v_h}{d} t \right) \quad (45)$$

Si des électrons ont été collectés, mais que les trous dérivent toujours ($t_e < t < t_h$) :

$$Q(t) = q_0 \left(\frac{x}{d} + \frac{v_h}{d} t \right) \quad (46)$$

Si des trous ont été collectés, mais que les électrons dérivent toujours ($t_h < t < t_e$) :

$$Q(t) = q_0 \left(\frac{v_e}{d} t + \frac{(d-x)}{d} \right) \quad (47)$$

Une fois que les trous et les électrons ont été collectés ($t > t_h$ et $t > t_e$) :

$$Q(t) = q_0 \quad (48)$$

La figure (17) montre plusieurs formes d'impulsion de ce type pour différentes valeurs de x . La variation du temps de montée effectif des impulsions est d'environ un facteur 2 puisque la vitesse de dérive des électrons est environ deux fois celle des trous.

b. Géométrie coaxiale

La forme de l'impulsion générée par la collection d'électrons et de trous dans un détecteur coaxial partage bon nombre des caractéristiques générales dérivées ci-dessus pour un type planaire, mais les parties linéaires du front d'attaque de l'impulsion prennent une courbure en raison de la géométrie différente.

Pour le cas simple des configurations Ge(Li) où aucune charge fixe nette n'est présente dans la région intrinsèque, le résultat est alors donné par :

$$Q(t) = \frac{q_0}{\ln(r_2/r_1)} \left[\ln \left(1 + \frac{v_e t}{r_0} \right) - \ln \left(1 - \frac{v_h t}{r_0} \right) \right] \quad (49)$$

L'équation (49) n'est valable que pour le premier domaine temporel, pendant lequel les trous et les électrons dérivent. Comme dans le cas précédent pour la géométrie planaire, les premier et deuxième termes entre crochets deviennent des constantes [égales à $\ln(r_2/r_0)$ et $\ln(r_1/r_0)$] lorsque les électrons ou les trous sont collectés, et un changement de pente se produit dans la forme d'onde aux instants correspondants. Une fois les deux collectés, $Q(t) = q_0$. Des tracés de cette forme d'impulsion

pour divers rayons d'interaction supposés sont donnés sur la figure (18). Ici, la variation du temps de montée effectif peut être encore plus grande que dans le cas planaire.

Les changements de forme de l'impulsion observés à partir d'un détecteur coaxial en fonction de la position radiale de l'interaction illustrée sur la figure (18) peuvent être exploités à certaines buts.

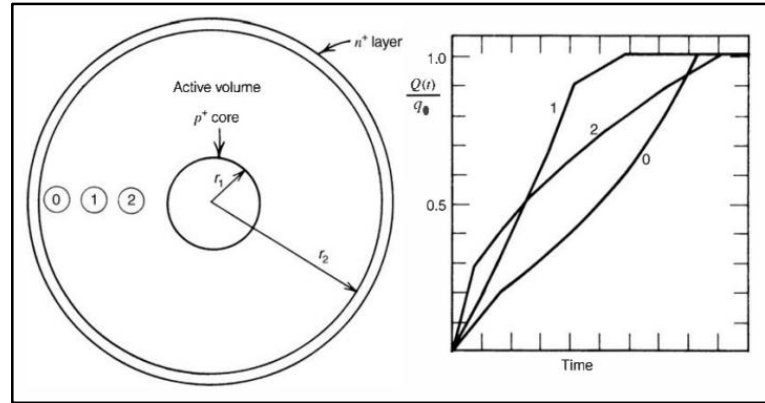


Fig. 18 Front montant de l'impulsion de sortie pour un détecteur à dérive coaxiale. Trois points d'interaction différents sont indiqués par 0, 1 et 2.

3.3.3 Effets du piégeage et du dépiégeage

Le piégeage des charges lors de leur dérive et leur éventuelle libération depuis des pièges peu profonds (dépiégeage) ont un effet significatif sur la forme attendue de l'impulsion. Pour simplifier la discussion suivante, nous supposons que toutes les paires électron-trou sont créées près d'une limite du volume actif. L'impulsion du signal correspond alors au mouvement uniquement de l'espèce porteuse collectée à la frontière opposée. Si l'on suppose que les porteurs se forment au contact p^+ d'un détecteur planaire, nous pouvons définir $x = d$ dans l'équation (45), qui se réduit alors à :

$$Q(t) = \begin{cases} q_0 \frac{t}{t_e} & (t \leq t_e) \\ q_0 & (t > t_e) \end{cases} \quad (50)$$

et la simple montée d'impulsion linéaire correspond à la dérive des électrons à travers le volume actif. Si nous supposons maintenant qu'une concentration uniforme de pièges à électrons existe dans cette région, une partie de la charge pourrait être perdue, au moins temporairement, de l'impulsion du signal. Si aucun dépiégeage n'a lieu, la perte est permanente et l'impulsion résultante prend la forme :

$$Q(t) = \begin{cases} \frac{q_0 \tau_T}{t_e} (1 - e^{-t/\tau_T}) & (t \leq t_e) \\ \frac{q_0 \tau_T}{t_e} (1 - e^{-t_e/\tau_T}) & (t > t_e) \end{cases} \quad (51)$$

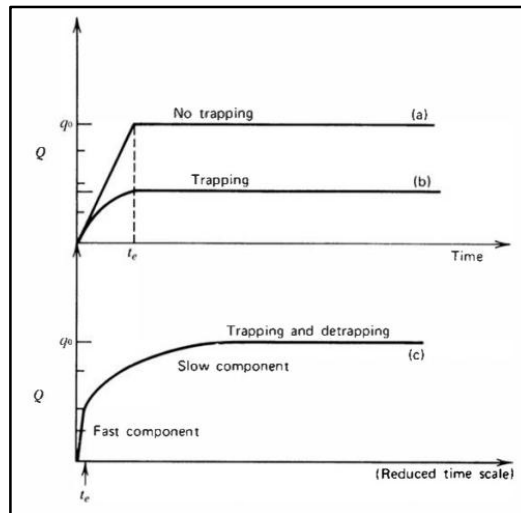


Fig. 19 Schéma illustrant le front montant de l'impulsion pour les interactions dans un détecteur planaire à une extrémité du volume actif : (a) pas de piégeage, (b) piégeage permanent, et (échelle de temps réduite), (c) piégeage avec dépiégeage lent.

Le piégeage peut ainsi conduire à une réduction de l'amplitude attendue de l'impulsion ou, si un dépiégeage se produit, à un ralentissement d'une partie du front montant de l'impulsion. Ce dernier peut avoir un effet significatif sur les propriétés temporelles des détecteurs au germanium dans lesquels le piégeage est important. De plus, le piégeage peut entraîner une détérioration de la résolution énergétique en raison de la quantité variable de charge perdue par impulsion. Le dépiégeage n'est utile à cet égard que s'il se produit sur une courte période de temps par rapport aux temps de formation des impulsions dans l'électronique ultérieure.

Dans les détecteurs coaxiaux, la géométrie cylindrique complique encore la prédiction de l'amplitude d'impulsion attendue en présence de piégeage.

3.3.3 Impulsions lentes ou défectueuses

Une fraction des impulsions observées à partir de certains détecteurs au germanium ont des temps de montée nettement plus longs que prévu par l'analyse ci-dessus. De plus, ces impulsions lentes sont également « défectueuses » en amplitude et peuvent contribuer à une perte de résolution énergétique ou à des queues de faible énergie sur les pics enregistrés dans le spectre de hauteur d'impulsion. Il existe des preuves que ces impulsions défectueuses correspondent à des événements qui prennent naissance près des bords du volume actif ou dans d'autres régions où le champ électrique local est inférieur à la normale. La collecte des charges est alors ralentie et les pertes dues au piégeage sont plus sévère que pour les impulsions normales. Ces effets sont amplifiés dans les détecteurs plus anciens présentant une forte concentration de défauts ou de centres de piégeage, souvent le résultat de dommages causés par les radiations dus à une exposition prolongée du détecteur à des neutrons rapides.

Étant donné que les impulsions à montée lente sont également déficientes en amplitude, des gains peuvent parfois être réalisés à la fois en termes de résolution énergétique et de rapport crête/Compton dans les spectres de rayons gamma en rejeter de telles impulsions en fonction de la détection de leurs longs temps de montée dans un circuit externe. Cependant, cette approche réduira également l'efficacité de la détection si une fraction substantielle de toutes les impulsions est impliquée. Une technique de correction de la hauteur d'impulsion plus élaborée peut à la place être utilisée pour éviter cette perte d'efficacité.

4. Spectroscopie Gamma avec des détecteurs au Germanium

Pour la mesure des énergies gamma supérieures à plusieurs centaines de keV, il existe deux catégories de détecteurs d'importance majeure : les scintillateurs inorganiques et les détecteurs à semi-conducteurs au germanium. En gardant à l'esprit que certains des matériaux de scintillation offrent une résolution énergétique améliorée et d'autres propriétés, à des fins de comparaison, nous utiliserons NaI (Tl) pour caractériser les scintillateurs inorganiques en raison de sa disponibilité économique en grandes tailles et de son utilisation actuelle généralisée. Bien qu'il existe de nombreux autres facteurs potentiels, le choix dans une application donnée repose le plus souvent sur un compromis entre l'efficacité du comptage et la résolution énergétique. Les scintillateurs à l'iodure de sodium présentent l'avantage d'être disponibles en grandes tailles, ce qui, associé à la densité élevée du matériau, peut entraîner des probabilités d'interaction très élevées pour les rayons gamma. Le numéro atomique relativement élevé de l'iode garantit également qu'une grande partie de toutes les interactions se traduira par une absorption complète de l'énergie des rayons gamma, et donc de la photo-fraction (ou fraction des événements située sous le pic de pleine énergie dans le spectre de hauteur d'impulsion) sera également relativement élevée.

La résolution énergétique des scintillateurs est mauvaise. Les spectres comparatifs illustrent la nette supériorité des détecteurs au germanium dans les situations où de nombreuses énergies de rayons gamma rapprochées doivent être séparées. Les bons systèmes au germanium auront une résolution énergétique typique de quelques dixièmes de pour cent, contre 5 à 10 % pour l'iodure de sodium. Cette résolution énergétique améliorée n'est cependant pas sans prix. Les plus petites tailles disponibles et le numéro atomique plus faible du germanium entraînent des efficacités photo-pic d'un ordre de grandeur inférieur dans les cas typiques. De plus, les photo-fractions sont faibles et les continus dans les spectres de hauteur d'impulsion posent donc davantage de problèmes.

Ce dernier inconvénient est quelque peu compensé par la résolution énergétique supérieure du germanium. Non seulement une bonne résolution aide à séparer les pics rapprochés, mais elle facilite également la détection de sources faibles d'énergies discrètes lorsqu'elles sont superposées sur un

large continuum. Les détecteurs ayant une efficacité égale donneront des zones égales sous le pic, mais ceux ayant une bonne résolution énergétique produiront un pic étroit mais haut qui peut alors s'élever au-dessus du bruit statistique du continuum.

Les détecteurs au germanium sont clairement préférés pour l'analyse de spectres gamma complexes impliquant de nombreux pics. Le choix est moins évident lorsque seules quelques énergies de rayons gamma très espacées sont impliquées, en particulier si l'objectif principal est une mesure de leur intensité plutôt qu'une détermination précise de l'énergie. Ensuite, la plus grande efficacité, la plus grande photo-fraction et le moindre coût des scintillateurs pourraient bien faire pencher la balance en leur faveur.

Annexe III

C - Autres détecteurs à semi-conducteurs

1. Détecteur de Silicium à dérive au Lithium

En utilisant du silicium de la plus haute pureté actuellement disponible, la profondeur d'appauvrissement pouvant être obtenue par polarisation inverse d'une diode au silicium conventionnelle est limitée à 1 à 2 mm. Si des détecteurs de silicium plus épais sont nécessaires, une approche différente doit être utilisée pour fabriquer le dispositif. Le processus de dérive du lithium peut plutôt être appliqué pour créer une région de silicium compensé ou « intrinsèque » dans laquelle les concentrations d'impuretés accepteurs et donneurs sont exactement équilibrées sur des épaisseurs allant jusqu'à 5 à 10 mm. Lorsqu'il est pourvu de électrodes non injectantes (souvent appelées contacts de blocage), cette région sert alors de volume actif à un détecteur. De tels détecteurs sont appelés détecteurs au silicium à dérive de lithium et reçoivent la désignation Si(Li). Leur épaisseur n'est limitée que par la distance sur laquelle la dérive du lithium peut être réalisée avec succès.

Le processus de compensation des impuretés par dérive lente du lithium dans un cristal semi-conducteur a été initialement développé pour produire des détecteurs à partir de germanium ou de silicium.

Pendant deux décennies, les détecteurs à dérive de germanium [ou Ge(Li)] représentaient la seule configuration dans laquelle un volume actif relativement important pouvait être fourni en germanium. Plus tard, des détecteurs de volume équivalent sont devenus disponibles dans la configuration au germanium de haute pureté (ou HPGe), et la dérive du lithium dans le germanium n'est plus populaire. Cependant, le procédé reste important dans le silicium car la pureté de son matériau disponible n'a pas égalé celle du germanium.

Le numéro atomique plus faible du silicium ($Z = 14$) par rapport à celui du germanium ($Z = 32$) signifie que pour les énergies typiques des rayons gamma, sa section efficace photoélectrique est inférieure d'un facteur 50 environ. L'efficacité maximale des rayons gamma à pleine énergie pour les détecteurs au silicium est donc très faible et, par conséquent, ils ne sont pas largement utilisés dans la spectroscopie gamma générale.

Il existe cependant deux domaines d'application principaux dans lesquels le numéro atomique inférieur du silicium n'est pas un obstacle mais une aide. Un domaine concerne la détection de rayons gamma ou rayons X de très faible énergie, où la probabilité d'absorption photoélectrique peut être raisonnablement élevée, même pour des détecteurs au silicium de quelques millimètres d'épaisseur. Le deuxième domaine d'application commun est la détection et la spectroscopie de particules bêta ou d'autres électrons incidents de l'extérieur.

Les propriétés semi-conductrices des deux matériaux favorisent dans une certaine mesure le silicium par rapport au germanium. Sa bande interdite plus grande garantit que le courant de fuite généré thermiquement, à une température donnée, sera plus petit par unité de volume dans le silicium que dans le germanium. L'énergie requise pour créer une paire électron-trou et le facteur Fano ne sont pas très différents dans les deux matériaux, de sorte que la contribution statistique à la résolution énergétique est d'environ 25 % égale. Si le piégeage de charge est négligeable, la résolution énergétique utilisant des composants électroniques équivalents peut être meilleure dans le silicium en raison de son courant de fuite plus faible.

1.1 Le processus de dérive ionique

Dans le silicium comme dans le germanium, le matériau ayant la pureté disponible la plus élevée a tendance à être de type p, dans lequel les meilleurs procédés de raffinage ont laissé une prédominance d'impuretés accepteurs. Des atomes donneurs doivent donc être ajoutés au matériau pour obtenir la compensation souhaitée. Les métaux alcalins tels que le lithium, le sodium et le potassium ont tendance à former des donneurs interstitiels dans les cristaux de silicium ou de germanium. Les atomes donneurs ionisés qui sont créés lorsque l'électron donné est excité dans la bande de conduction sont suffisamment mobiles à des températures élevées pour pouvoir dériver sur des périodes de plusieurs jours sous l'influence d'un champ électrique puissant. Parmi les exemples mentionnés ci-dessus, seul le lithium peut être introduit dans le silicium ou le germanium en concentration suffisante pour servir de dopant compensateur pratique. L'ion lithium chargé positivement possède un noyau semblable à l'hélium qui lui confère un petit rayon favorisant sa diffusion à travers d'autres solides.

Les ions lithium peuvent être introduits dans des cristaux de silicium et de germanium pour produire des détecteurs. La mobilité des ions lithium est beaucoup plus grande dans le germanium et reste suffisamment élevée à température ambiante pour permettre une redistribution indésirable du lithium à partir de la situation compensée obtenue pendant la dérive. Le profil du lithium doit donc être préservé immédiatement après la dérive dans le germanium en réduisant drastiquement la température des cristaux, typiquement à celle de l'azote liquide (77 K). Dans le silicium, la mobilité des ions est suffisamment faible à température ambiante pour permettre le stockage temporaire des détecteurs au silicium à dérive de lithium sans refroidissement.

1.2 La configuration PIN

La plupart des détecteurs Si(Li) sont fabriqués selon une géométrie planaire dans laquelle le lithium dérive de la surface plane d'une épaisse plaquette de silicium. Bien que certaines

configurations coaxiales à grand volume aient été produites, les applications courantes des détecteurs Si(Li) à la mesure des rayons X mous ou des électrons bénéficient davantage de la grande surface d'entrée plate qu'offre la configuration planaire.

Une fois le processus de dérive terminé, le détecteur résultant a la configuration simplifiée illustrée sur la figure (20). L'excès de lithium à la surface à partir de laquelle la dérive a commencé sert à convertir cette surface en une couche n^+ pouvant être utilisée comme contact électrique. La région p non compensée du côté opposé reçoit souvent un revêtement métallique qui agit comme un contact ohmique.

La durée de vie des porteurs de charge créés dans la région compensée (souvent simplement appelée région intrinsèque ou i) peut être considérablement supérieure au temps nécessaire pour les collecter à l'une ou l'autre des limites, et de bonnes propriétés de collecte de charges peuvent donc en résulter. Il est cependant nécessaire de collecter les charges rapidement et des tensions importantes sont normalement appliquées pour garantir que peu d'électrons ou de trous soient perdus avant que la collecte ne soit terminée. Les tensions de polarisation typiques sont de 500 à 4 000 V dans des régions intrinsèques de 5 à 10 mm d'épaisseur.

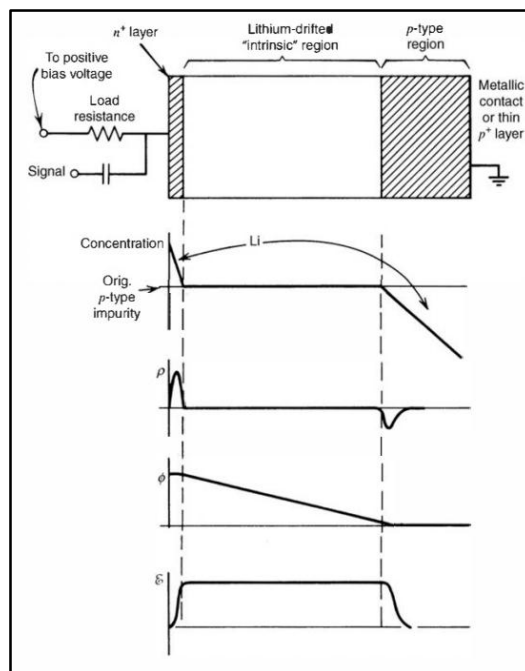


Fig. 20 Configuration de base d'un détecteur de jonction pin à dérive de lithium, les profils correspondants pour la concentration d'impuretés, la densité de charge ρ , le potentiel électrique ϕ et le champ électrique ξ .

2. Matériaux semi-conducteurs autres que le Silicium ou le Germanium

Contrairement aux semi-conducteurs constitués d'un seul élément chimique représenté par le silicium (Si) et le germanium (Ge), c'est un semi-conducteur composé (AgCl) qui a été utilisé pour la

première fois pour démontrer que les signaux induits par les rayonnements pouvaient être extraits d'un dispositif semi-conducteur. Les semi-conducteurs composés ont connu un développement lent au fil des années, largement inhibé par des problèmes de croissance cristalline impliquant des impuretés électriquement actives et des défauts natifs. Les problèmes de manipulation ont également été considérables, dans la mesure où certains semi-conducteurs composés ne possèdent pas les propriétés matérielles robustes requises pour résister à la technologie courante des processus de semi-conducteurs. Dans de nombreux cas, les détecteurs à semi-conducteurs composés ont montré une instabilité temporelle, souvent appelée polarisation.

Les semi-conducteurs composés les plus étudiés pour les mesures de rayonnement sont CdTe, $\text{Cd}_{1-x}\text{Zn}_x\text{Te}$ et HgI_2 , leurs valeurs de bande interdite sont suffisamment grandes pour permettre un fonctionnement à température ambiante pour les applications typiques de rayons gamma. Cependant, un bon nombre des applications les plus prometteuses impliquent des mesures aux rayons X dans lesquelles l'énergie déposée et le signal de charge induit qui en résulte sont relativement faibles. Dans de tels cas, un refroidissement modeste peut réduire le courant de fuite et améliorer ainsi la résolution énergétique de ces détecteurs. Les détecteurs de haute température n'ont besoin que d'être refroidis à 0°C pour obtenir des courants de fuite inférieurs à 0,1 pA dans les cas typiques. Une réduction supplémentaire de la température peut réduire le niveau de bruit du circuit d'entrée sur le préamplificateur, mais toute réduction supplémentaire du courant de fuite déjà faible du détecteur n'a pas d'effet appréciable sur la résolution énergétique du détecteur. Dans le cas du CdTe, le courant de fuite est suffisamment élevé pour qu'un refroidissement supplémentaire à -40°C puisse produire des améliorations utiles dans leur résolution énergétique. La compacité est préservée grâce à l'utilisation de refroidisseurs Peltier électriques miniatures intégrés au boîtier du détecteur et dont la présence est pratiquement invisible pour l'utilisateur.

2.1 Détecteurs au tellure de cadmium (CdTe)

Le tellure de cadmium (CdTe) combine des numéros atomiques relativement élevés (48 pour le Cd et 52 pour le Te) avec une énergie de bande interdite suffisamment grande (1,52 eV) pour permettre un fonctionnement à température ambiante ou proche.

La probabilité d'absorption photoélectrique par unité de longueur de trajet est environ 4 à 5 fois plus élevée dans le CdTe que dans le Ge, et 100 fois plus grande que dans le Si pour les énergies typiques des rayons gamma.

Les applications de ce matériau impliquent donc le plus souvent des situations dans lesquelles une efficacité élevée de détection des rayons gamma par unité de volume est primordiale.

2.2 Détecteurs d'iodure mercurique (HgI₂)

L'iodure mercurique est un autre matériau qui a suscité une attention considérable en tant que détecteur potentiel des rayons X et des rayons gamma. HgI₂ est un matériau semi-conducteur avec une grande énergie de bande interdite de 2,13 eV, qui permet un fonctionnement à température ambiante tout en ne produisant qu'un très faible courant de fuite, généré thermiquement. La section efficace photoélectrique élevée conduit à une efficacité d'absorption des rayons gamma de basse énergie bien supérieure à celle du Ge, à certaines énergies jusqu'à 50 fois supérieures. À cet égard, HgI₂ est un détecteur de photons encore plus efficace que CdTe. Par exemple, 85 % des photons de 100 keV sont absorbés dans une épaisseur de 1 mm de HgI₂, alors que le même résultat nécessiterait 2,6 mm de CdTe et 10 mm de Ge.

2.3 Détecteurs Cd_{1-x} Zn_x Te (« CZT »)

Il a été constaté que le mélange de ZnTe dans CdTe donne des cristaux présentant des propriétés favorables en tant que détecteurs, notamment une résistivité significativement plus élevée. Le sigle « CZT » est généralement utilisé pour ce matériau. Le résultat est un composé ternaire de formulation Cd_{1-x} Zn_x Te, où la fraction de mélange x peut varier de 0,04 à 0,20. Les semi-conducteurs résultants ont des énergies de bande interdite à température ambiante de 1,53 eV et 1,64 eV, respectivement. L'efficacité d'absorption des rayons gamma est similaire à celle du CdTe, mais l'augmentation de la bande interdite réduit la concentration intrinsèque de porteurs libres et le courant de fuite de fonctionnement en dessous de celui généralement observé avec des appareils CdTe. Les mobilités mesurées sont de 1 350 cm²/V-s pour les électrons et de 120 cm²/V-s pour les trous, les caractéristiques de vitesse de dérive sont donc significativement différentes. Semblable au CdTe, la durée de vie mesurée des trous est beaucoup plus courte que la durée de vie des électrons, étant comprise entre 100 ns et plusieurs microsecondes pour les électrons et entre 50 et 300 ns pour les trous. Le tableau (01) résume les propriétés de quelques matériaux semi-conducteurs

Tab. 01 Propriétés des matériaux semi-conducteurs

Material	Z	Density (g/cm ³)	Bandgap (eV)	Ionization Energy (eV per $e-h$ pair)	Best Gamma-Ray Energy Resolution (FWHM)
Si (300 K)	14	2.33	1.12	3.61	
(77 K)			1.16	3.76	400 eV at 60 keV
(77 K)					550 eV at 122 keV
Ge (77 K)	32	5.33	0.72	2.98	400 eV at 122 keV
					900 eV at 662 keV
					1300 eV at 1332 keV
CdTe (300 K)	48/52	6.06	1.52	4.43	1.7 keV at 60 keV
					3.5 keV at 122 keV
HgI ₂ (300 K)	80/53	6.4	2.13	4.3	3.2 keV at 122 keV
					5.96 keV at 662 keV
Cd _{0.8} Zn _{0.2} Te (300 K)	48/30/52	6	1.64	5.0	11.6 keV at 662 keV

2.4 Autres détecteurs à semi-conducteurs

2.4.1 Arséniure de gallium (GaAs)

L'arséniure de gallium (GaAs) a été étudié comme détecteur de rayonnement depuis le début des années 1960 et a été le premier semi-conducteur composé fonctionnant à température ambiante qui a démontré une bonne résolution énergétique des rayons gamma. L'énergie de bande interdite à température ambiante du GaAs est de 1,43 eV et l'énergie d'ionisation moyenne est de 4,3 eV/e-h, ce qui indique qu'une résolution énergétique acceptable devrait être possible pour un fonctionnement à température ambiante. Les numéros atomiques de Ga (31) et d'As (33) sont proches de celui du Ge, de sorte que les interactions attendues des rayons gamma et l'efficacité de détection par unité de masse devraient être similaires à celles des détecteurs au germanium. Les premiers spectromètres GaAs expérimentaux étaient fabriqués à partir d'un matériau épitaxial en phase liquide de très haute pureté, mais des problèmes de coût excessif et de répétabilité ont empêché leur introduction en tant que détecteurs commercialement viables. Les tentatives visant à utiliser du matériau GaAs cultivé en masse pour les spectromètres à particules chargées et à rayons gamma n'ont connu qu'un succès modeste.

2.4.2 Iodure de plomb (PbI₂)

L'iodure de plomb (PbI₂) reste intéressant comme semi-conducteur alternatif pour la spectroscopie gamma. Les grands numéros atomiques du plomb (82) et de l'iode (53) indiquent que le matériau aura une très bonne interaction avec les rayons gamma et une très bonne efficacité de détection. Son énergie de bande interdite est de 2,55 eV, ce qui implique que le matériau peut fonctionner dans des environnements bien au-dessus de la température ambiante. Bien que le PbI₂ ait été étudié pour la première fois au début des années 1970 en tant que détecteur de rayonnement, des problèmes importants concernant son utilisation demeurent. Des améliorations récentes ont été réalisées, mais la faible mobilité des porteurs de charge, associée à un piégeage important des porteurs de charge, a continué à empêcher le PbI₂ de devenir un spectromètre à semi-conducteurs commercialement viable.

2.4.3 Bromure de thallium (TlBr)

Le bromure de thallium (TlBr) offre des propriétés intéressantes en tant que matériau de détection candidat et a suscité une attention considérable dans la communauté des chercheurs. Sa grande énergie de bande interdite de 2,68 eV permettrait un fonctionnement à température ambiante ou même au-dessus, bien qu'elle entraîne également une énergie relativement importante nécessaire pour créer une paire électron-trou (6,5 eV). Sa haute densité (7,56 g/cm³) et son nombre atomique

élémentaire élevé (81 et 35) favorisent une excellente efficacité de détection des rayons gamma, même dans des cristaux relativement minces.

L'utilisation de couches de silicium amorphe comme détecteurs de rayonnement a également suscité un certain intérêt. Le silicium sous forme amorphe offre le potentiel d'un coût bien inférieur et d'une surface plus grande par rapport aux monocristaux classiquement utilisés dans la fabrication des détecteurs au silicium. Les valeurs de mobilité et de durée de vie des électrons et des trous dans les matériaux amorphes sont inférieures de plusieurs ordres de grandeur à celles observées dans les cristaux de silicium. En conséquence, seules des couches très fines (jusqu'à 30 μm) ont été utilisées avec succès comme détecteurs, et encore uniquement avec une collecte de charges fractionnées. Les couches de ce type, bien que généralement inutiles en spectroscopie énergétique, peuvent fonctionner comme de simples compteurs de particules chargées pour des applications dans lesquelles la résolution énergétique n'est pas importante. Les impulsions des particules alpha ont été mesurées pour des couches aussi fines que 5 μm .

Une certaine utilisation du diamant comme matériau de détection a été réalisée dans des applications spécialisées. Le diamant est un matériau isolant avec une très grande bande interdite (5,6 eV) et peut être utilisé comme un simple compteur de conduction en appliquant des contacts électriques sur les faces opposées du cristal. Les détecteurs de ce type, bien qu'évidemment limités à de petites tailles, se caractérisent par un temps de réponse rapide de plusieurs nanosecondes et une stabilité supérieure à long terme. Ils peuvent fonctionner soit en mode impulsif conventionnel, soit en mode courant dans des champs de rayonnement élevés. Les mesures ont démontré une résolution énergétique de 22 keV pour la spectroscopie à température ambiante des particules alpha. En raison de leur large bande interdite, les détecteurs de diamant sont bien adaptés au fonctionnement à des températures élevées, et leur utilisation à 250-300°C a été rapportée.

Plusieurs autres semi-conducteurs ont été étudiés avec un succès limité, notamment InP, GaSe, et CdSe. Le tableau ci-dessous répertorie certaines propriétés de divers matériaux semi-conducteurs à des fins de comparaison. Aucun de ces matériaux n'a atteint le point de viabilité commerciale. La plupart sont limités à de très petites tailles et il existe souvent des problèmes persistants en matière de collecte efficace des charges.

Les travaux de développement se poursuivent cependant dans l'espoir de produire une alternative au Si et au Ge capable de fonctionner à température ambiante.

Des détecteurs intéressants ont également été fabriqués à partir de carbure de silicium (SiC), un autre matériau à bande interdite élevée (3,27 eV). Les valeurs Z sont trop faibles pour être intéressantes pour la spectroscopie des rayons gamma, mais des applications réussies pour d'autres mesures de rayonnement. Le matériau existe dans un certain nombre de poly-types liés aux

spécificités de la structure cristalline, et la forme 4H (désignée 4H-SiC) a été le choix le plus courant. L'une des propriétés frappantes de ces dispositifs est un courant de fuite extrêmement faible qui ouvre la possibilité d'une application à très faible bruit à la spectroscopie des rayons X mous sans nécessiter de refroidissement. Leur excellente résistance aux dommages causés par les radiations et leur opérabilité à des températures élevées ont conduit à d'autres applications dans des environnements à fort rayonnement ou dans la détection de neutrons. Le tableau (02) résume les propriétés de quelques matériaux semi-conducteurs composés.

Tab. 02 Certains matériaux semi-conducteurs composés

Material	Z	Density (g/cm ³)	Bandgap (eV)	Ionization Energy (eV per <i>e-h</i> pair)	Best Gamma-Ray Energy Resolution (FWHM)
GaAs	31/33	5.32	1.43	4.2	2.2–2.5 keV at 60 keV 2.6 keV at 122 keV
InP	49/15	4.79	1.35	4.2	—
CdSe	48/34	5.8	1.73	—	8.5 keV at 60 keV
PbI ₂	82/53	6.2	2.55	4.9	1.83 keV at 60 keV
GaSe	31/34	4.55	2.03	4.5	—
TlBr	81/35	7.56	2.68	6.5	—

3. Détecteurs d'avalanche

Le détecteur à semi-conducteur normal est l'analogue à l'état solide d'une chambre d'ionisation conventionnelle remplie de gaz. L'objectif dans les deux cas est simplement de collecter tous les porteurs de charges qui ont été amenés par le rayonnement incident dans le volume actif. Dans certaines conditions, il est également possible d'atteindre la multiplication de charge dans les détecteurs à l'état solide. Le dispositif résultant est alors l'analogue du compteur proportionnel et est appelé détecteur d'avalanche. Bien qu'ils ne soient pas aussi largement appliqués que les compteurs proportionnels à gaz, ils ont atteint un certain degré de popularité pour la détection de rayonnements à faible énergie, y compris les rayons X.

Le gain est réalisé dans le matériau semi-conducteur (normalement le silicium) en augmentant le champ électrique suffisamment élevé pour permettre aux électrons migrés de créer une ionisation secondaire pendant le processus de collecte. Les conceptions du détecteur d'avalanche doivent incorporer une géométrie de dispositif qui permet à la génération de champs élevés dans le volume du détecteur sans permettre au champ à la surface de devenir si grand que de créer une dégradation de la surface. Cet effet est parfois accompli en contournant la forme du détecteur comme illustré sur la figure (21).

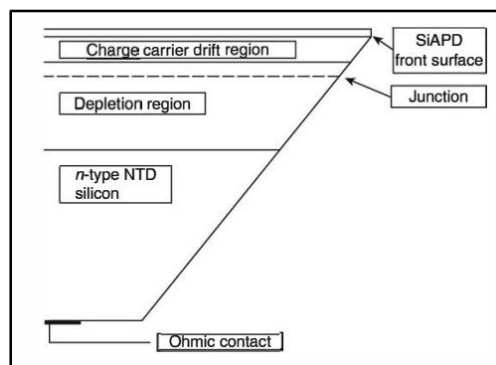


Fig. 21 Coupe transversale d'une diode à avalanche montrant le bord biseauté utilisé pour éviter la dégradation de la surface

Des gains allant jusqu'à plusieurs centaines de charges totales collectées sont possibles, mais il a été difficile d'obtenir une multiplication uniforme sur toute la fenêtre d'entrée de ces détecteurs. Le gain dépend également fortement de la tension et de la température appliquées, et un fonctionnement stable nécessite une attention particulière à ces facteurs.

Les détecteurs d'avalanches offrent des propriétés intéressantes, tant du point de vue de la détection des petits signaux que de la rapidité de réponse. Les champs très élevés garantissent que le temps de transit à travers la région de multiplication est très court et que les temps de montée des impulsions typiques sont de l'ordre de 3 ns. La résolution temporelle typique est inférieure à la nanoseconde et peut être inférieure à 0,1 ns dans des circonstances favorables. Cette combinaison de propriétés permet l'utilisation de détecteurs d'avalanche présentant de bonnes caractéristiques signal/arrière-plan, même pour la détection de rayonnements de très faible énergie ou sous conditions dans lesquelles la température du détecteur est élevée. Par exemple, les rayons X de 1,5 keV peuvent être détectés avec une efficacité intrinsèque de 30 à 40 % à des températures comprises entre 85 et 100°C tout en maintenant des niveaux de fond très faibles.

Les détecteurs d'avalanches ont trouvé de nombreuses applications dans la détection des rayons X de faible énergie. D'autres applications ont impliqué l'utilisation dans des moniteurs portables de rayons gamma et dans la détection de particules bêta de faible énergie provenant du tritium.

4. Détecteurs photoconducteurs

Un type spécialisé de détecteur à semi-conducteurs s'est révélé utile dans la mesure de courtes rafales de rayonnement intense. Dans ces conditions, de nombreuses particules ou photons interagissent dans le volume actif dans le temps nécessaire pour dériver les charges résultantes vers les électrodes. Un détecteur photoconducteur est constitué d'un échantillon de matériau semi-conducteur équipé de deux contacts d'injection (ou contacts « ohmiques ») sur des surfaces opposées. Lorsqu'une tension est appliquée, un courant mesurable circule (le courant d'équilibre) qui est

déterminé par les concentrations de porteurs libres dans le matériau semi-conducteur. La concentration en porteurs sera augmentée en irradiant le dispositif avec une impulsion (ou un flux continu) de rayonnement ionisant, ce qui entraînera une augmentation de la conductivité du matériau. Le photo-courant intégré (l'augmentation par rapport au courant d'équilibre), ou la charge photoconductrice, est proportionnel à l'énergie déposée par l'impulsion de rayonnement.

Le courant de fuite est très faible dans les matériaux à large bande interdite ou à haute résistivité, mais ils nécessitent également une attention particulière lors de la fabrication de véritables contacts d'injection. Sans injection de charges efficace, les conditions nécessaires au gain photoconducteur ne sont plus satisfaites.

5. Détecteurs à semi-conducteurs sensibles à la position

Les détecteurs dans lesquels la position d'interaction du rayonnement incident est détectée ainsi que son énergie ont des applications dans un certain nombre de domaines différents. Les deux principaux types de détecteurs utilisés pour la détection de position des particules chargées sont les dispositifs remplis de gaz et les détecteurs à diodes semi-conductrices au silicium ou au germanium. Ces derniers types sont parfois préférés en raison de leur compacité et de leur faible tension de polarisation par rapport aux détecteurs à gaz. Les détecteurs à semi-conducteurs, en raison de leur plus grand pouvoir d'arrêt, sont également mieux adaptés aux applications impliquant des rayonnements à longue portée.

5.1 Division de charge résistive

Dans sa forme la plus simple, le détecteur semi-conducteur sensible à la position est constitué d'une bande unidimensionnelle de silicium ou de germanium pour laquelle un contact est réalisé pour avoir une résistance série significative. Comme le montre la figure (22), une extrémité de ce contact résistif est mise à la terre tandis que l'autre mène à un amplificateur pour la dérivation du signal de position. Le contact résistif agit comme un diviseur de charge et la quantité de charge délivrée à l'amplificateur de signal de position est proportionnelle à x/L , où x est la distance de l'interaction depuis l'extrémité mise à la terre et L est la longueur de la bande. Un deuxième signal (le signal E) proportionnel à la charge totale déposée dans le détecteur est dérivé de l'électrode avant conductrice normale. Si le signal de position est divisé par le signal E , une impulsion est produite dont l'amplitude reflète la position de l'interaction, indépendamment de la quantité de charge réellement déposée.

Dans certaines conditions, la partie non appauvrie d'une barrière superficielle normale en Silicium peut être utilisée comme couche arrière résistive. Plus typiquement, le contact résistif est formé par évaporation métallique ou par techniques d'implantation ionique. Les couches implantées

par des ions sont généralement supérieures car elles peuvent être rendues uniformes en termes de résistance.

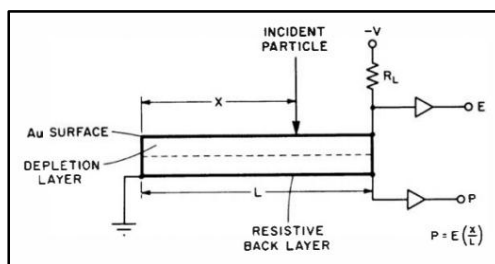


Fig. 22 Configuration de base d'un détecteur à semi-conducteur sensible à la position. Le signal de position (P) est obtenu par division de charge résistive de la couche arrière. Le signal E est proportionnel à l'énergie déposée par la particule dans le volume actif du détecteur

5.2 Détecteurs micro-ruban

Une autre méthode de détection de position consiste à subdiviser l'une des électrodes en un certain nombre de segments ou bandes indépendants. Étant donné que les paires électron-trou créées dans le volume du détecteur se déplaceront le long des lignes de champ jusqu'au segment d'électrode correspondant, un signal fort sera dérivé uniquement des segments qui ont collecté des porteurs de charge appréciables. Si des particules à courte portée sont impliquées, un seul signal de ce type sera généré. Pour les particules à plus longue portée, le « centre de gravité » de la position de la trace peut être déterminé par interpolation des signaux provenant de segments proches. Un détecteur à micro-ruban de silicium est un détecteur dans lequel une série d'électrodes à bandes parallèles étroites ont été fabriquées sur une surface à l'aide de techniques d'implantation ionique et/ou de photolithographie. Pour une résolution spatiale fine, de tels détecteurs ont été produits avec des largeurs de bande aussi petites que 10 μm . Pour éviter la multiplicité de canaux électroniques indépendants qui seraient nécessaires pour mesurer indépendamment le signal de chaque bande, des systèmes ont été développés pour trouver la bande la plus proche de l'interaction en utilisant une division de charge unidimensionnelle le long d'une ligne qui relie toutes les bandes. En utilisant cette approche, on peut obtenir une résolution spatiale approchant les 10 μm dans une dimension pour les particules alpha incidentes. Le traitement double face de la tranche de silicium peut produire des configurations comme illustré sur la figure (23) dans lesquelles des bandes d'électrodes sont prévues dans une orientation orthogonale sur les côtés opposés de la tranche.

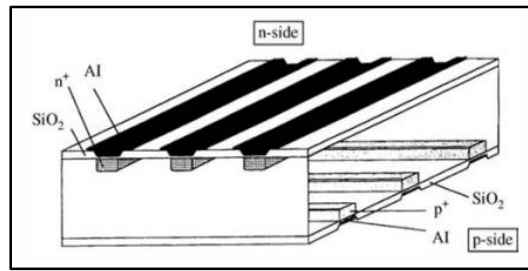


Fig. 23 Détecteur à bande de silicium double face

Les détecteurs à bande de silicium ont trouvé une application généralisée dans la poursuite de particules à haute énergie dans des expériences basées sur des accélérateurs et ont également été appliqués à l'imagerie des rayons X. Il existe également des applications pour lesquelles des détecteurs de silicium à bandes orthogonales de plus grande épaisseur et/ou surface sont nécessaires, par exemple l'imagerie des rayons gamma via le processus de diffusion Compton. Ensuite, la plus grande épaisseur offerte par le type de détecteur Si(Li) décrit plus haut dans ce chapitre est préférée, et les dimensions des bandes sont désormais généralement à l'échelle des millimètres plutôt que des micromètres.

5.3 Détecteurs aux tampons ou aux pixels

L'approche par « force brute » pour obtenir des informations de position bidimensionnelles à partir d'un détecteur au silicium unilatéral consiste à fabriquer l'électrode supérieure sous la forme d'un motif en damier composé de petits éléments individuels électriquement isolés les uns des autres. Lorsque les dimensions des électrodes individuelles sont d'un millimètre ou plus, ce type d'appareil est généralement appelé détecteur à tampon. Pour les dimensions d'électrode inférieures à un millimètre, la terminologie courante est : détecteur aux pixels. Une connexion électrique doit être établie pour chaque électrode individuelle et des canaux de lecture électronique séparés doivent être prévus pour chacune. Un avantage de cette approche est que la petite taille de chaque électrode individuelle entraîne une capacité et un courant de fuite, relativement faibles, et ainsi le bruit électronique est considérablement réduit par rapport à celui observé avec des détecteurs micro-ruban de dimensions équivalentes.

Fournir des contacts électriques séparés à chaque pixel présente un défi important. Pour les détecteurs de petite surface dans lesquels les pixels ou les plages sont suffisamment grands, des techniques de liaison par fil peuvent être utilisées pour assurer une connexion électrique de chaque pixel au bord de la tranche. Alternativement, certaines constructions à double couche métallique permettent le contact avec les pixels intérieurs. Cependant, dans l'approche la plus courante, illustrée

sur la figure (24), une puce de détection de pixels est connectée à une puce de lecture séparée en utilisant une liaison par soudure à puce retournée ou des liaisons à bosses d'indium.

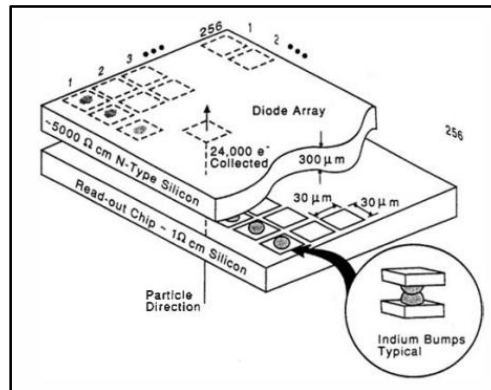


Fig. 24 Détecteur aux pixels « hybride » composé d'un détecteur séparé et de puces électroniques connectées via des liaisons à choc en indium

Les détecteurs à pixels ou à pastilles ont généralement des zones actives limitées à quelques centimètres carrés. Des zones de détection plus grandes peuvent être obtenues en assemblant des modules individuels en réseaux, mais au prix d'une complexité croissante dans ce qui est déjà un dispositif complexe.

5.4 Détecteurs de dérive à semi-conducteurs

Il a été démontré que le temps de dérive des porteurs de charge peut être utilisé pour déduire la position de leur formation dans un détecteur au silicium. Dans un détecteur à dérive de semi-conducteur, une configuration d'électrode unique est utilisée comme illustre la figure (25). Des jonctions semi-conductrices sont formées sur les deux faces d'une tranche de grande surface, et chacune est polarisée en inverse jusqu'à ce que le détecteur soit complètement épuisé.

La figure (25) montre un dispositif dit à dérive linéaire dans lequel des bandes parallèles sont utilisées pour façonner le champ et créer le gradient de potentiel nécessaire pour conduire les électrons vers l'anode collectrice. Dans ce cas, l'anode est segmentée pour permettre la détermination de la deuxième coordonnée de position dans la dimension parallèle aux bandes.

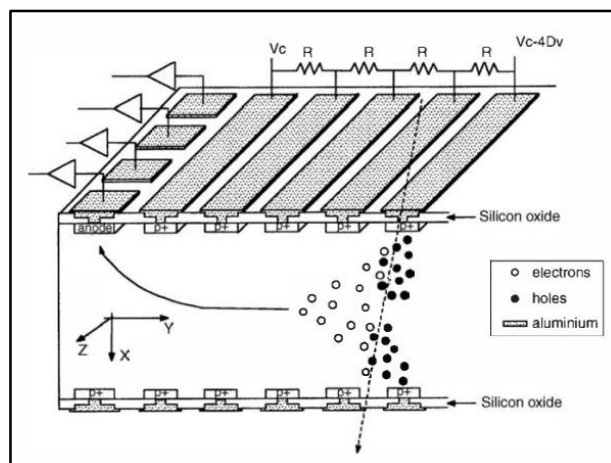


Fig. 25 Détecteur à dérive linéaire dans lequel des bandes parallèles sont utilisées pour façonner le champ

La figure (26) montre une géométrie de dérive cylindrique dans laquelle l'anode collectrice est située au centre d'une série d'anneaux circulaires qui servent à façonner le champ électrique. Dans ce cas, tous les électrons d'ionisation sont collectés sur une seule anode de petite surface qui reste très petite pour minimiser sa capacité.

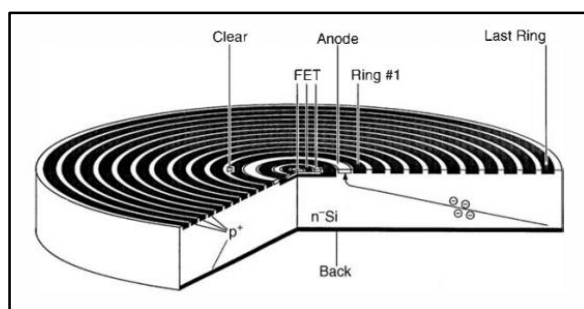


Fig. 26 Coupe transversale d'un détecteur de dérive cylindrique

La configuration du détecteur de dérive présente un avantage supplémentaire qui peut être exploité pour améliorer la résolution énergétique. Étant donné que les électrons peuvent dériver sur de grandes distances et être collectés sur une électrode de très petite taille, la capacité du détecteur peut être bien inférieure à celle d'une diode semi-conductrice équivalente de conception conventionnelle. La capacité du détecteur est toujours un facteur important pour déterminer le niveau de bruit d'un système spectroscopique, et une faible capacité constitue un avantage significatif. Les détecteurs de dérive ont donc retenu une attention considérable en tant que spectromètres à rayons X.

La petite capacité des détecteurs de dérive ou de pixels constitue également un avantage dans les applications à taux de comptage élevé. Étant donné que le temps de mise en forme optimal pour minimiser le bruit électronique se déplace vers des valeurs inférieures à mesure que la capacité du

détecteur diminue, des temps de mise en forme plus courts sont possibles, ce qui conduit à une capacité de taux de comptage plus élevée pour un degré d'empilement donné.

5.5 Appareils à couplage de charge

Les dispositifs à couplage de charge (CCD) ont été introduits pour la première fois au début des années 1970 en tant que dispositifs à microcircuits en silicium capables d'enregistrer des images en lumière visible. Leurs applications se sont développées rapidement et les CCD se retrouvent dans une large gamme de composants optiques et de caméras. Ces applications commerciales restent la force économique dominante derrière le développement des CCD. Cependant, un sous-ensemble plus restreint de dispositifs dotés de propriétés similaires, souvent appelés CCD scientifiques, est apparu dans les années 1990 comme capteur extrêmement utile pour la détection et l'imagerie des rayonnements. Ils ont été largement utilisés dans le suivi ou l'imagerie de particules ionisantes minimales de haute énergie. Les CCD sont également devenus une alternative un peu plus complexe mais viable aux autres détecteurs au silicium pour la spectroscopie de routine aux rayons X, en particulier aux basses énergies.

Un diagramme schématisant une architecture CCD simple est présenté sur la figure (27). Ce sont normalement des structures de silicium produites par de nombreuses techniques de fabrication microélectronique standard utilisées pour les circuits intégrés.

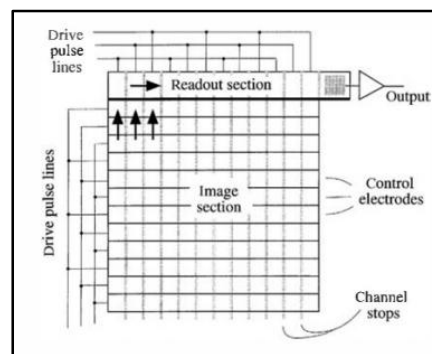


Fig. 27 Aménagement de la surface d'un CCD

Le CCD et le détecteur de dérive décrits précédemment partagent de nombreuses propriétés communes. Les deux assurent le transport des électrons créés par une particule ionisante sur de longues distances parallèles à la surface de la plaquette de silicium. Le CCD y parvient par un défilement séquentiel synchronisé de rangées de pixels vers la section de lecture, tandis que le détecteur de dérive pilote les charges en continu. Ce dernier processus est intrinsèquement plus rapide et il est intéressant de combiner les caractéristiques des deux concepts pour surmonter certaines des limitations des CCD, principalement leur lente lecture.

5.6 Imageurs à rayons X à semi-conducteurs à couche épaisse

Les détecteurs d'imagerie numérique à rayons X constituent une technologie de détection de plus en plus importante pour une utilisation en imagerie médicale et en inspection industrielle. Des matériaux semi-conducteurs à couches épaisses (généralement appelés couches photoconductrices) sont utilisés pour de tels dispositifs et forment une classe de matériaux détecteurs de rayonnement distincts des plaquettes semi-conductrices cristallines plus conventionnelles.

En particulier, la capacité de fabriquer des panneaux détecteurs de rayonnement en utilisant la technologie des semi-conducteurs à couches épaisses est très intéressante pour les capteurs d'imagerie pour lesquels des détecteurs « à écran plat » à grande surface active sont souhaitables. Dans la technologie conventionnelle des détecteurs à semi-conducteurs, la zone active du détecteur est normalement limitée soit par la disponibilité de tranches semi-conductrices de taille suffisante, soit par les contraintes de connexion de la couche semi-conductrice à l'électronique de lecture requise. En revanche, l'utilisation d'une couche semi-conductrice à couche épaisse, généralement cultivée directement sur un type de circuit de lecture de charge à base de silicium, peut produire des détecteurs d'imagerie de grande surface dans lesquels la zone active n'est plus limitée par les semi-conducteurs cristallins classiques.

L'une des approches les plus réussies utilisées pour fabriquer des détecteurs d'imagerie pixellisés sur de grandes surfaces a été les réseaux de transistors à couches minces (TFT) en silicium amorphe (a-Si). Le réseau TFT a-Si est recouvert d'une couche de conversion appropriée dans laquelle les interactions de rayonnement se produisent, qui est soit un phosphore optique ou un scintillateur, soit une couche semi-conductrice pour la production directe de charges (Fig. 28). Dans le cas d'une couche de convertisseur optique, les interactions de rayonnement produisent des photons optiques dans le spectre visible qui sont détectés par le TFT, qui fait office de photo-capteur.

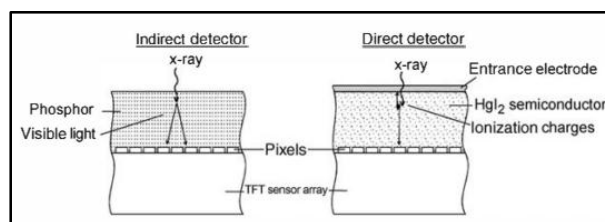


Fig. 26 Coupe transversale schématique des détecteurs d'imagerie TFT, montrant la détection indirecte via des photons optiques et la détection directe de charge à l'aide d'une couche semi-conductrice

Alternativement, la production directe de charge dans une couche semi-conductrice présente plusieurs avantages potentiels par rapport à la conversion optique, à savoir un rapport signal/bruit plus élevé et une résolution spatiale améliorée. Le type le plus important commercialement de

détecteurs d'imagerie TFT à conversion directe utilise du sélénium amorphe (a-Se) comme couche de conversion de charge directe.

Ces appareils ont été développés commercialement comme détecteurs de radiographie numérique dans les années 1980 et 1990 et ont démontré une excellente qualité d'image. Des conceptions alternatives de TFT ont également été développées en utilisant des matériaux autres que l'a-Si pour la matrice du transistor, par exemple le silicium polycristallin ou le Séléniure de cadmium polycristallin, qui peuvent offrir des performances de transistor améliorées en raison d'une mobilité plus élevée des porteurs.

En raison du numéro atomique du Se ($Z = 34$) nettement plus élevé que celui du silicium, le rendement pour les photons de 50 keV est proche de 100 %.

Le principal avantage potentiel des réseaux TFT à détection directe a-Se, par rapport aux détecteurs TFT à scintillateur ou à revêtement phosphore, est la résolution spatiale améliorée, qui résulte de la conversion directe des rayons X en charge électronique.

Une variété de matériaux semi-conducteurs à couches épaisses a été étudiée en tant que couches photoconductrices à rayons X qui constituent des alternatives potentielles à l'a-Se. L'utilisation de semi-conducteurs à Z plus élevé offre une plus grande efficacité quantique à une énergie de photons de rayons X plus élevée et des améliorations potentielles de la qualité de l'image et une réduction de la dose au patient. Suite aux développements parallèles des semi-conducteurs composés monocristallins, les deux matériaux à couches épaisses les plus avancés pour les détecteurs sont le tellure de cadmium (CdTe) et l'iodure de mercure (HgI_2). Ces matériaux, qui ont des valeurs Z efficaces moyennes de 50 et 80, respectivement, peuvent être produits sous forme de films épais polycristallins par diverses techniques de croissance telles que le dépôt thermique ou la sublimation en espace clos. Les semi-conducteurs composés offrent également un gain significatif en sensibilité de charge par rapport à l'a-Se, en raison de l'énergie de création de paires électron-trou plus faible (environ 5 eV/e-h paire contre 50 eV/e-h paire dans a-Se).

Outre le HgI_2 , plusieurs autres matériaux semi-conducteurs composés ont été étudiés en tant que photoconducteurs à couche épaisse pour les capteurs d'imagerie TFT et CMOS. Bien qu'aucun ne soit aussi développé commercialement que HgI_2 , des prototypes de dispositifs d'imagerie ont été démontrés utilisant du CdTe polycristallin, de l'iodure de plomb (PbI_2) et de l'oxyde de plomb (PbO). Il a été démontré que les matrices d'imagerie prototypes PbI_2 TFT affichent des performances raisonnables ; cependant, le décalage d'image et la sensibilité aux rayons X inférieure ont tendance à être pires que pour les dispositifs HgI_2 correspondants, et des améliorations supplémentaires sont encore nécessaires pour la qualité du matériau semi-conducteur. Des résultats impressionnants ont été obtenus avec des détecteurs d'imagerie TFT à base de PbO , avec des zones actives allant jusqu'à

18 cm x 20 cm. Des couches de PbO polycristallin de haute qualité d'une épaisseur allant jusqu'à 300 μm ont été fabriquées par évaporation thermique à basse température.

Les matériaux intermédiaires $Z \sim 50$ tels que CdTe et CdZnTe ont également été étudiés comme photoconducteurs à couches épaisses, aussi bien sous forme poly-cristalline qu'épitaxiale. Cependant, malgré le succès considérable des détecteurs de rayonnement monocristallins CdTe et CdZnTe pour les applications de spectroscopie, le développement actuel des détecteurs d'imagerie à couches épaisses CdTe/CdZnTe est encore limité au laboratoire de recherche. Des prototypes de détecteurs d'imagerie poly-cristallins CdZnTe ont été rapportés. Cependant, en raison de la température élevée requise pour le processus de dépôt par sublimation en espace clos de CdZnTe, le matériau a été développé sur un substrat d'alumine puis lié au réseau TFT. D'autres groupes ont caractérisé le CdZnTe polycristallin déposé sur des substrats par évaporation thermique, mais la température de dépôt relativement élevée de 350°C a empêché le dépôt direct sur des matrices TFT.

Le tableau 03 résume les propriétés électriques et de transport de charge de certains matériaux détecteurs poly-cristallins à Z élevé, avec un produit typique de mobilité électronique-durée de vie pour le HgI_2 polycristallin de $5 \times 10^{-5} \text{ cm}^2/\text{V}$.

Tab. 03 Comparaison des propriétés électriques et de transport de charge des matériaux de détecteurs poly-cristallins à conversion directe

Material	Effective Electron-Hole Pair Creation Energy (eV)	Atomic Number	Density (g/cm^3)	Resistivity ($\Omega \text{ cm}$)	Electron Mobility-Lifetime Product (cm^2/V)	Thickness to Stop 63% of 60 keV X-Rays (μm)
a-Se	~ 50	34	4.3	$> 10^{12}$	$> 10^{-9}$	990
CdZnTe	4.5	48, 30, 52	~ 6	$> 10^{10}$	$> 10^{-5}$ $> 10^{-3} SC$	260
PbI_2	5	82, 53	6.2	$> 10^{11}$	$> 10^{-8}$	250
PbO	—	82, 8	9.5	—	$> 10^{-7}$	
HgI_2	4.2	80, 53	6.4	$> 10^{12}$	$> 10^{-5}$ $> 10^{-3} SC$	250

Bibliographies

T. E. Valentine, *Annals of Nuclear Energy* 28, 191 (2001).

R. D. Evans, *The Atomic Nucleus*, Krieger, New York, 1982.

"Neutron and Gamma-Ray Fluence-to-Dose Factors," American National Standard ANSI/ANS-6.1.1-1991.

J. M. Harrer and J. G. Beckerley, *Nuclear Power Reactor Instrumentation Systems Handbook*, Vol. 1, Chap. 5, TID-25952-PI (1973).

J. Sharpe, *Nuclear Radiation Detectors*, 2nd ed., Methuen and Co., London, 1964.

C.A.N. Conde, "Gas Proportional Scintillation Counters for XRay Spectrometry" in *X-ray Spectrometry: Recent Technological Advances*, K. Tsuji, J. Injuk, and R. Van Grieken (eds.), John Wiley & Sons, 2004.

W. J. Price, *Nuclear Radiation Detection*, 2nd ed., Chap. 5, McGraw-Hill, New York, 1964.

J.B.Birks, *The Theory and Practice of Scintillation Counting*, Pergamon Press, Oxford, 1964

S. M. Shafroth (ed.), *Scintillation Spectroscopy of Gamma Radiation*, Gordon & Breach, London, 1967.

G. Lutz, *Semiconductor Radiation Detectors*, Springer-Verlag, Berlin, 1999.

R. A. I. Bell, "Tables for Calibration of Radiation Detectors," Australian National University Report ANU-P/606 (1974).

Abstract

This course is intended mainly for undergraduate students (physics specialty), its objective is to teach the basic elements of radiation detectors while focusing on the study of the general properties of these systems. Since the operating principle of radiation detectors is based on the properties of the interaction of incident radiation with the material constituting the detector, it was necessary to begin with a reminder of the elementary notions of the physics of radiation-matter interactions. Likewise, the course as a whole is designed to develop a good understanding of the physical mechanisms of detectors which are the basis of their applications in the field of radiation detection.

Understanding how a detector works therefore requires knowing these interaction processes. This is the reason why this course brings together the fundamental notions concerning radiation and their sources as well as their interactions with matter in addition to the different families of detectors and their general properties.

Keywords: Interaction radiation-matter, radiation detectors, gaseous detectors, solid detectors, scintillation detectors.

الملخص

هذا المقرر مخصص بشكل رئيسي لطلبة المرحلة الاولى الجامعية (تخصص فيزياء)، وهدفه هو دراسة العناصر الأساسية لأجهزة الكشف عن الإشعاع مع التركيز على دراسة الخصائص العامة لهذه الأنظمة. بما أن مبدأ تشغيل أجهزة الكشف عن الإشعاع يعتمد على خصائص تفاعل الإشعاعات مع المادة المكونة للكاشف، فقد كان من الضروري البدء بتذكير بالمفاهيم الأولية لفيزياء تفاعلات المادة مع الإشعاعات. وعلى العموم لقد تم تصميم المقرر في مجمله لتطوير فهم جيد للآليات الفيزيائية لأجهزة الكشف التي تشكل أساس تطبيقاتها في مجال الكشف عن الإشعاع. ولذلك فإن فهم كيفية عمل الكاشف يتطلب معرفة عمليات التفاعل هذه. ولهذا السبب يجمع هذا المقرر المفاهيم الأساسية المتعلقة بالإشعاع ومصادره وتفاعلاته مع المادة بالإضافة إلى عائلات الكواشف المختلفة وخصائصها العامة.

الكلمات المفتاحية: تفاعل الإشعاعات مع المادة ، كواشف الإشعاع ، الكواشف الغازية ، الكواشف الصلبة ، الكواشف الوامضة.