



UNIVERSITÀ DEGLI STUDI DI CATANIA

FACOLTÀ DI INGEGNERIA

CORSO DI LAUREA IN INGEGNERIA MICROELETTRONICA

GIOVANNI RICCARDO GUARDO

**ANALISI E CARATTERIZZAZIONE DI
FOTOSENSORI AL SILICIO PER SINGOLO
FOTONE PER IMPIEGO IN RIVELATORI
DI RADIAZIONE NUCLEARE**

TESI DI LAUREA

Relatore:

Chiar.mo Prof. G. Palumbo

Correlatore:

Dott. P. Finocchiaro

ANNO ACCADEMICO 2009-2010

*“Posso deporre la mia anima, per
poi riprenderla una seconda volta”*

Physiologus

Indice

Capitolo I:

Oggetto della tesi	4
---------------------------------	----------

Capitolo II:

Introduzione	6
---------------------------	----------

Capitolo III:

Interazione dei raggi gamma con la materia	16
---	-----------

Capitolo IV:

Scintillatori	21
----------------------------	-----------

Capitolo V:

Silicon PhotoMultiplier	27
--------------------------------------	-----------

Capitolo VI:

Strumentazione	35
-----------------------------	-----------

Capitolo VII:

Misure di risoluzione temporale	42
--	-----------

Capitolo VIII:

Misure in coincidenza	57
------------------------------------	-----------

Capitolo IX:	
Misure di PDE	66
Capitolo X:	
Analisi di dark count	76
Capitolo XI:	
Conclusioni	87
Bibliografia	88

Capitolo I:

Oggetto della tesi

Il lavoro di tesi, come suggerisce il titolo, è volto allo studio analitico e alla caratterizzazione di fotomoltiplicatori al silicio (SiPM) di diverse case costruttrici al fine di determinarne l' idoneità per l' utilizzo nel progetto **TOPEM** (TOF-PET MRI).

Questo progetto, sviluppato in collaborazione tra INFN (Istituto Nazionale di Fisica Nucleare) e ISS (Istituto Superiore di Sanità), si propone la realizzazione di una sonda per la diagnosi del cancro alla prostata basata su *Time of Flight PET*, ma che integri al suo interno anche le funzionalità della *MRI*.

Sotto queste premesse l' utilizzo di dispositivi microelettronici offre notevoli vantaggi, ad esempio, in termini di costi e ingombro del dispositivo finale o di integrazione di PET e MRI, la quale è favorita dalla naturale immunità ai campi magnetici che presentano i SiPM.

La tecnica di diagnostica qui suggerita presenta, oltre ai suddetti, altri vantaggi rispetto alla tradizionale PET, nonché altre problematiche, che però approfondiremo nei capitoli seguenti.

Capitolo II:

Introduzione

La tomografia a emissione di positroni (o PET dall'inglese Positron Emission Tomography) è una tecnica di medicina nucleare e di diagnostica medica utilizzata per la produzione di mappe dei processi funzionali interni al corpo, usata estensivamente in oncologia clinica e nelle ricerche cardiologiche e neurologiche.

La procedura inizia con l'iniezione per via endovenosa o l'inalazione di gas (rari casi), nel soggetto da esaminare, di un isotopo tracciante con tempo di dimezzamento breve legato chimicamente a una molecola attiva a livello metabolico.

Dopo un tempo di attesa (da qualche minuto fino a due ore, in base al tipo d'isotopo), durante il quale la molecola metabolicamente attiva, raggiunge una determinata concentrazione all'interno dei tessuti organici da analizzare, il soggetto occupa posto nello scanner.



Figura 2.1: Scanner PET.

L'isotopo a questo punto decade, emettendo un positrone che, annichilandosi con un elettrone, producendo una coppia di fotoni gamma emessi in direzioni opposte fra loro, che sono rilevati da un **rivelatore di radiazione nucleare** a scintillatori situato nel dispositivo di scansione, dove creano un lampo luminoso, che è rilevato attraverso degli opportuni fotosensori

Il processo viene reiterato ruotando una griglia radiale di rivelatori permettendo così una ricostruzione tridimensionale dell'attività nucleare nel corpo.

Punto cruciale della tecnica è la rilevazione simultanea di coppie di eventi: i fotoni che non raggiungono il rivelatore in coppia entro un intervallo di pochi nanosecondi, non sono presi in considerazione.

Dalla misurazione della posizione in cui i fotoni colpiscono il rivelatore, si può ricostruire la posizione del corpo da cui sono stati emessi, permettendo la determinazione dell'attività o dell'utilizzo chimico all'interno degli organi investigati, mediante una mappa della densità dell'isotopo nel corpo.

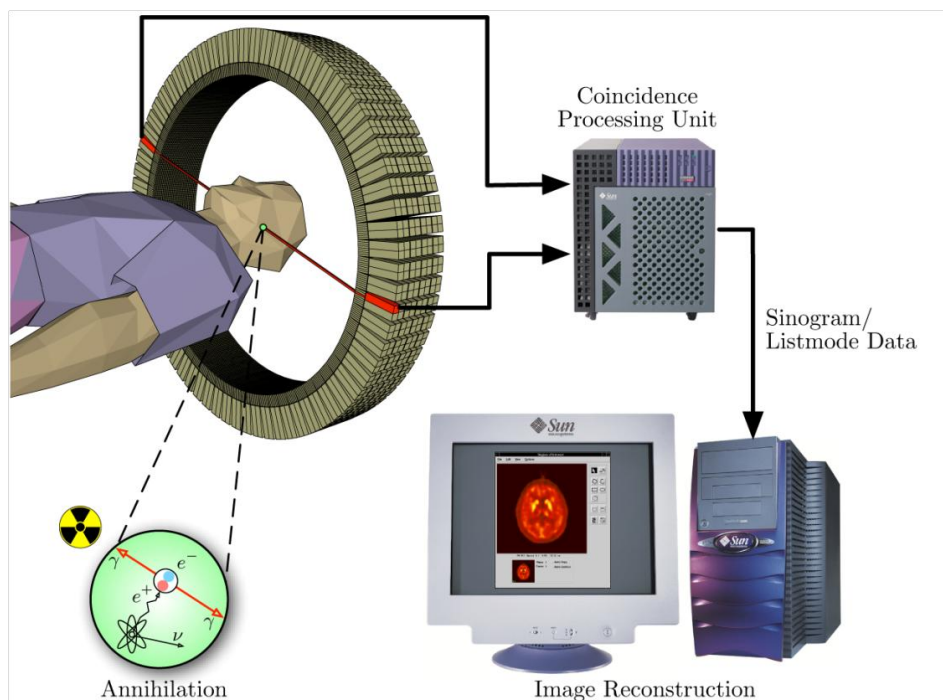


Figura 2.2: Schema di funzionamento della PET.

Apriamo una piccola parentesi sui rivelatori che entrano in gioco nel funzionamento.

Nella fisica sperimentale, un **rivelatore di particelle** è uno strumento elettromeccanico usato per rivelare, tracciare e identificare particelle ionizzanti, come quelle prodotte per esempio

da un decadimento nucleare, dalla radiazione cosmica o dalle interazioni in un acceleratore di particelle.

In particolare, quando lo scopo principale è la misura della radiazione, si preferisce usare il termine **rivelatore di radiazione**.

Molti tipi di rivelatori sfruttano la creazione di coppie dovuta al passaggio della radiazione, per esempio coppie elettrone-ione in un rivelatore a gas (come ad esempio il ben noto *contatore Geiger*), o elettrone-lacuna in un rivelatore a semiconduttore.

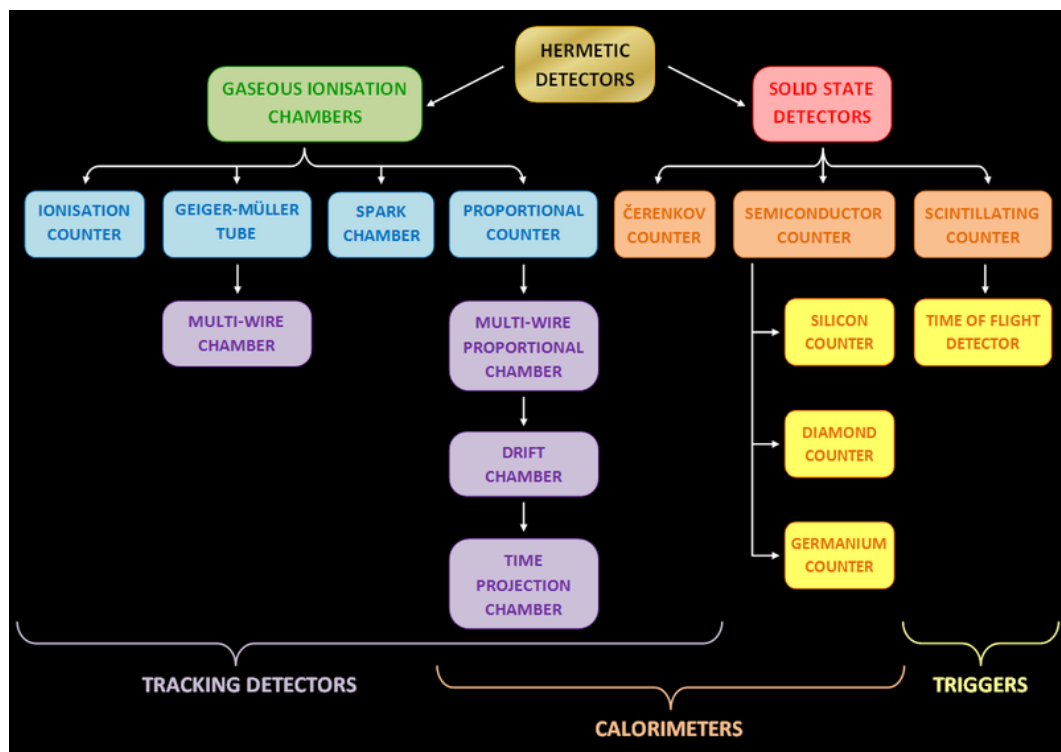


Figura 2.3: Schema riassuntivo dei principali rivelatori di particelle.

Il rivelatore sviluppato e caratterizzato durante il lavoro di tesi si colloca nella parte destra dello schema proposto precedentemente ed è costituito da uno *scintillatore*, un materiale capace di emettere

impulsi di luce quando viene attraversato da fotoni di alta energia, abbinato a dei *fotomoltiplicatori*, sensori capaci di rilevare ed amplificare tali impulsi di luce (I principi di funzionamento di questi due oggetti verranno illustrati dettagliatamente nei capitoli successivi).

Al giorno d'oggi il tumore alla prostata è il cancro più comune negli uomini residenti nei paesi occidentali, nonché la seconda causa di morte per cancro al mondo.

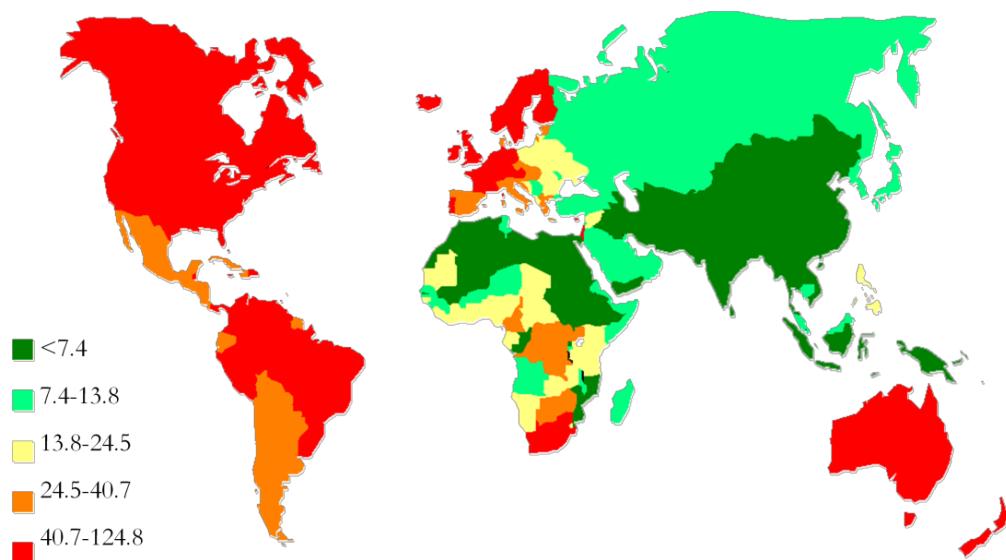


Figura 2.4: Incidenza del tumore su 100,000 individui.

I metodi di diagnostica sono molteplici, ma nessuno finora soddisfa pienamente in termini di *sensibilità* (capacità di individuare il possibile tumore) e di *specificità* (capacità di corretta diagnosi, cioè di discriminare il tumore da altre possibili patologie).

La tecnica di diagnosi più comune un'analisi del tasso di PSA (antigene prostatico specifico) nel sangue del soggetto, che ha un'ottima sensibilità, ma una pessima specificità in quanto, anche per valori superiori ai 10 ng/ml critici, l'incidenza del tumore è solo di un caso su due.

Per questi motivi il conteggio del PSA è usata solo come diagnosi preliminare seguita da successivi controlli.

Un'altra tecnica di diagnosi, che come principio di funzionamento è molto simile alla PET ma usa raggi X anziché raggi gamma, è la TAC, la quale ha una specificità piuttosto alta, ma pecca in sensibilità (specialmente per lesioni interne alla prostata o estensioni extracapsulari).

Infine abbiamo la *MRI (magnetic resonance imaging)*, che presenta una sensibilità mediamente tra il 75% e il 90%, ma può arrivare fino al 97% per lesioni note o valori inferiori al range tipico se le lesioni sono sconosciute e inferiori ai 5 mm.

Purtroppo la specificità è di poco superiore al 50% a causa dell'elevato numero di falsi positivi che si hanno nella mappatura.

Il responso definitivo si può avere solamente a seguito di una biopsia alla cieca che potrebbe causare lesioni involontarie, quindi bisognerebbe considerare la possibilità di un drastico cambiamento nell'approccio verso la diagnosi del cancro alla prostata.

Gli svantaggi nell'uso della PET tradizionale sono legati alla bassa risoluzione spaziale ($6 \div 12$ mm) e alla distanza dei sensori

dall'organo in esame, in quanto la radioattività esterna alla prostata causa una notevole perdita di contrasto nell'immagine, dato che una grossa concentrazione di radioisotopo è presente nella vescica.

I moderni sistemi permettono una migliore risoluzione temporale nella discriminazione del tempo di arrivo della coppia di fotoni, quindi è stata sviluppata una variante di questa tecnica di diagnostica, la PET a Tempo di Volo (o TOF-PET dall'inglese Time of Flight PET), nella quale si determina con precisione l'istante di arrivo di ciascun fotone della coppia, localizzando così il punto di origine dell'evento di annichilazione entro qualche centimetro e consentendo inoltre l'eliminazione del fondo rumoroso (come quello dato dalla vescica ad esempio).

In termini analitici si ha un numero di elementi di volume che contribuiscono al rumore che è dato solo da quelli associati all'informazione TOF, quindi un SNR notevolmente ridotto.

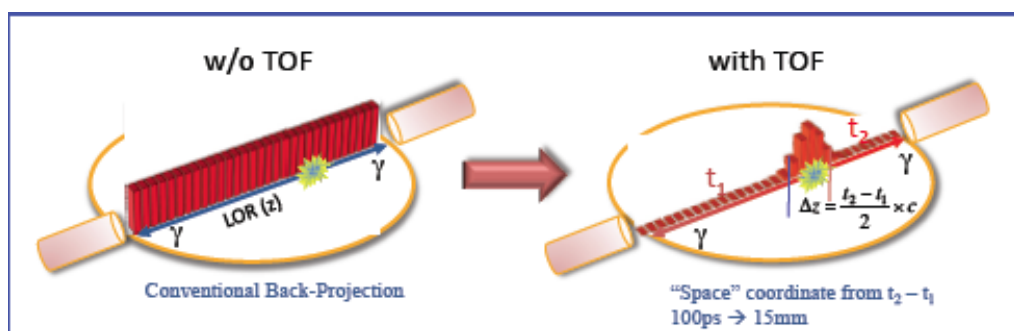


Figura 2.5: Miglioramento introdotto dalla tecnica TOF.

Per dare qualche numero basti pensare che la maggior parte degli scanner PET che non fanno uso della tecnica suddetta sia vincolata a una risoluzione temporale di circa 2 nsec, mentre con primi sistemi commercializzati con TOF-PET si ottengano risoluzioni intorno ai 600 psec.

Il concept di TOPEM prevede l'interfacciamento in coincidenza tra un rivelatore PET esterno e una sonda prostatica PET interna, con l'eventuale ausilio di alcuni sensori esterni ad alta risoluzione che migliorino l'imaging dei linfonodi superficiali.

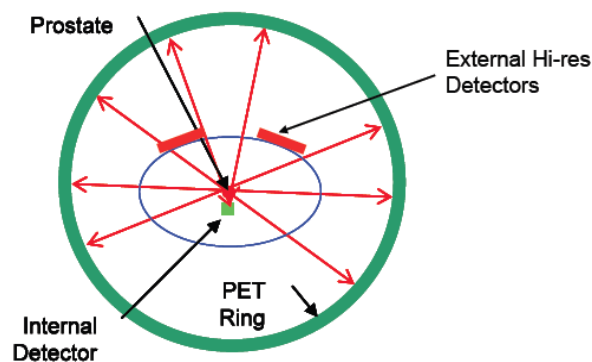


Figura 2.6: Schema strutturale del progetto.

Il rinnovato interesse da parte del settore della medicina nucleare in questa tecnica sta spingendo verso successivi miglioramenti della risoluzione temporale grazie all'ottimizzazione di tutti i parametri in essa coinvolti, quali il cristallo scintillatore, il fotosensore e l'elettronica di processamento del segnale.

Abbiamo pocanzi accennato all'importanza di una buona risoluzione temporale, poiché essa si traduce direttamente in una buona risoluzione spaziale, più approfonditamente ciò è espresso dalla banale legge:

$$\Delta x = \Delta t \, c$$

Dove c è la velocità della luce nel mezzo, espressa da:

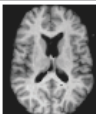
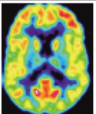
$$c = \frac{c_0}{n}$$

Con c_0 velocità della luce nel vuoto e n indice di rifrazione del mezzo.

Con un indice di rifrazione che mediamente per gli scintillatori è intorno a due e una risoluzione temporale di 600 ps si ottengono dunque 9 cm di risoluzione spaziale.

Come già premesso, TOPEM si muove nell'ottica di integrare al meglio i vantaggi di PET e MRI in un unico strumento.

Complementary nature of MRI & PET

Parameter	MRI 	PET 
Anatomical Detail	Excellent	Poor
Spatial Resolution	Excellent	Compromised
Clinical Penetration	Excellent	Limited
Sensitivity	Poor	Excellent
Molecular imaging	Limited	Excellent

Hence: The Sum of PET and MRI should be excellent and even better

MRI + PET << MRI-PET

Figura 2.7: Confronto tra PET e MRI.

Per ottenere ciò però bisogna considerare vari fattori indesiderati che entrano in gioco:

- I sensori PET non devono né essere realizzati con materiali magnetici, né emettere alle frequenze di MR, al fine di evitare distorsioni dell'immagine di MR.
- Come corollario del precedente punto, i sensori PET devono essere insensibili ai campi elettromagnetici (la trasmissione RF della MRI può produrre falsi eventi PET).
- I materiali per la schermatura PET devono essere MRI-compatibili e, in maniera duale, i materiali per la MRI devono attenuare il meno possibile i raggi gamma.
- Bisogna tener conto dei gradienti di campo e delle correnti indotte causate dalla MRI, che possono produrre riscaldamento e rumore nei sensori.

L'uso dei SiPM, come discuteremo più avanti, ben si presta a soddisfare alcuni dei suddetti punti, come quelli riguardanti l'insensibilità elettromagnetica ad esempio.

La necessità di una sonda che funzioni dall'interno, inoltre, ben si sposa con l'utilizzo di microsensori elettronici che hanno il vantaggio di avere, oltre ad un ingombro minimo, una tensione di alimentazione molto contenuta e un costo molto basso rispetto alle alternative possibili (tubi fotomoltiplicatori ad esempio).

Capitolo III:

Interazione dei raggi gamma con la materia

Quando andiamo a fare delle misure per rivelazione di radiazione nucleare bisogna innanzitutto capire che cosa ci si aspetta in termini di energia, in particolare nella reazione di annichilazione tra positrone ed elettrone, per la conservazione dell'energia e della quantità di moto, sono creati due fotoni aventi ciascuno un'energia pari a quella a riposo dell'elettrone o del positrone, ossia 511 keV.

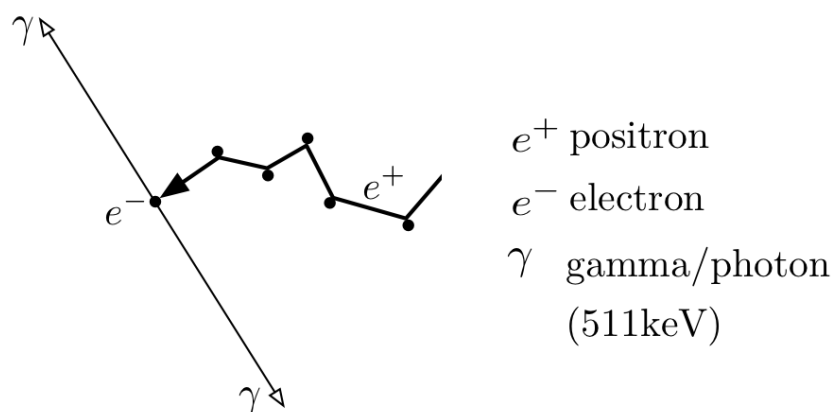


Figura 3.1: Processo di annichilazione con creazione della coppia di gamma.

Tali fotoni hanno delle frequenze che ricadono nell'intervallo dei raggi gamma poiché:

$$\nu = \frac{E}{h} \cong 10^{20} \text{Hz}$$

Inoltre, poiché il sistema possiede inizialmente una quantità di moto totale pari a zero, i gamma sono emessi in direzioni opposte.

Tra i tanti meccanismi d'interazione dei raggi gamma con la materia, però, soltanto tre giocano un importante ruolo nelle misure di radiazione, e ciascuno di essi è caratterizzato dalla cessione (parziale o totale) di energia da un fotone a un elettrone.

Il primo fenomeno d'interazione che esaminiamo, che è allo stesso tempo quello che ci interessa principalmente per ragioni che spiegheremo più avanti, è l'**effetto (o assorbimento) fotoelettrico**, in cui il raggio gamma scompare cedendo interamente la sua energia a un fotoelettrone, il quale è espulso dall'atomo con un'energia pari a

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

Dove E_b rappresenta l'energia di legame del fotoelettrone con la sua posizione originaria.

Per radiazioni gamma di basse energie (qualche centinaio di keV) il processo fotoelettrico è dominante, poiché il fotoelettrone porta con sé quasi tutta l'energia del fotone iniziale, e ciò è esaltato quando l'interazione è con materiali con alto numero atomico Z , proprietà

che ritornerà utile più avanti, quando discuteremo dei cristalli scintillatori.

Nell'**effetto Compton** il raggio gamma interagisce elasticamente con un elettrone trasferendo parte della sua energia allo stesso, l'elettrone così urtato è espulso, mentre il fotone subisce uno scatter con angolo φ .

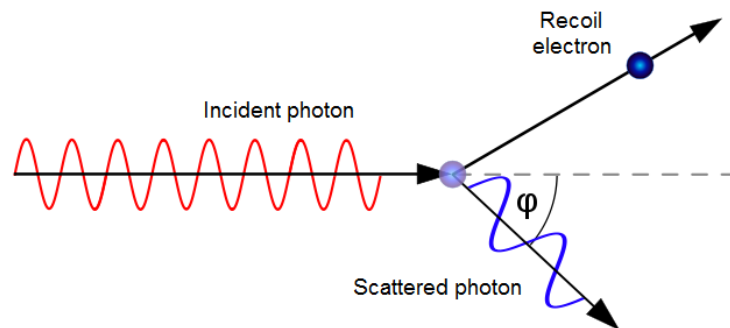


Figura 3.2: Rappresentazione del fenomeno di Compton Scattering.

Il fotone avrà un'energia minore rispetto a prima, in particolare applicando sempre la conservazione dell'energia e del momento si trova:

$$h \nu' = \frac{h \nu}{1 + \frac{h \nu}{m_0 c^2} (1 - \cos \varphi)}$$

dove $m_0 c^2$ è l'energia associata alla massa a riposo dell'elettrone (511 keV appunto).

Questo si traduce direttamente in un aumento della lunghezza d'onda detto **shift Compton**:

$$\Delta\lambda = \frac{h}{m_0 c} (1 - \cos\varphi)$$

È interessante osservare che, nel caso di piccoli angoli di diffusione, il fenomeno è pressoché ininfluenza.

Se il raggio gamma supera un'energia pari al doppio di quella dell'elettrone a riposo (1.02 MeV), la **produzione di coppie** è allora possibile; in questo terzo e ultimo meccanismo il raggio gamma scompare cedendo la sua energia per creare una coppia positrone-elettrone.

C'è da aggiungere però che la probabilità d'innescare di questo processo è effettivamente tangibile per radiazioni gamma con energia di almeno un ordine di grandezza superiore degli 1.02 MeV suddetti, quindi quest'ultimo effetto, poiché lavoreremo con gamma di basse energie, non sarà preso più in considerazione.

Riassumiamo quanto detto sinora esaminando un grafico che rappresenta la curva in cui le probabilità τ (assorbimento fotoelettrico) e σ (scattering Compton) si eguagliano, al variare dell'energia della radiazione e per diversi numeri atomici Z .

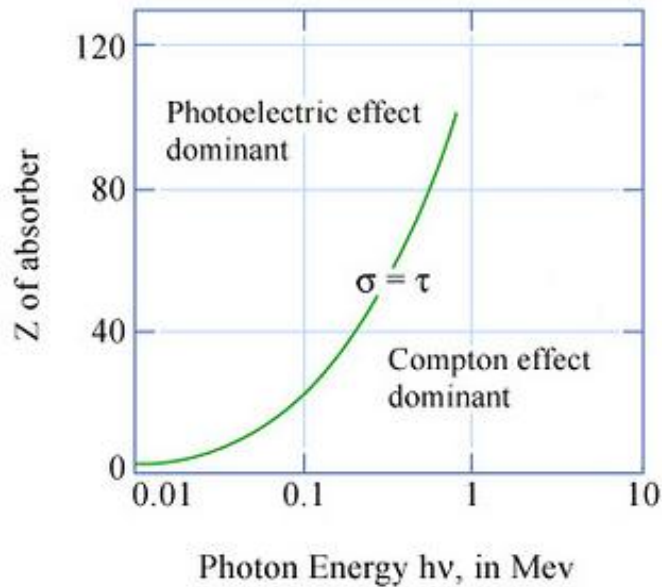


Figura 3.3: Grafico delle probabilità Compton e Fotoelettrica in funzione dell'energia e del numero atomico del mezzo.

È immediato che, come predetto, alti valori di Z e radiazioni di bassa energia accentuino l'assorbimento fotoelettrico, ma quello che emerge in seconda analisi è anche una diversa dipendenza da Z per le due probabilità, in particolare:

$$\begin{cases} \tau \propto Z^n \\ \sigma \propto Z \end{cases}$$

Dove n è un valore che varia tra 4 e 5.

Nel nostro caso, per quanto riguarda l'energia siamo vincolati al valore rilasciato dalla reazione di annichilazione, possiamo dunque lavorare, ma sempre entro certi limiti, su Z , ossia sul materiale del cristallo scintillatore.

Capitolo IV:

Scintillatori

Il processo di scintillazione è uno dei metodi più utili per la rilevazione e la spettroscopia di una grande varietà di radiazioni.

Lo scintillatore ideale dovrebbe possedere le seguenti proprietà:

1. Dovrebbe convertire l'energia cinetica delle particelle cariche in luce visibile con un'alta efficienza di scintillazione.
2. Questa conversione dovrebbe essere lineare, ossia la quantità di luce dovrebbe essere proporzionale all'energia depositata, entro una dinamica più ampia possibile.
3. Il mezzo dovrebbe essere trasparente alla lunghezza d'onda che emette per una raccolta ottimale della luce.
4. Il tempo di decadimento della luminescenza indotta dovrebbe essere breve, cosicché sia possibile generare segnali a impulsi veloci.
5. L'indice di rifrazione dovrebbe essere vicino a quello del materiale del fotorivelatore (vetro o resina a seconda se parliamo di fototubi o SiPM).

Nessun materiale soddisfa contemporaneamente tutti questi criteri, quindi la scelta di un particolare scintillatore è sempre un compromesso tra questi e altri fattori.

Gli scintillatori si distinguono fondamentalmente in organici e inorganici, i primi tendenzialmente hanno migliore linearità e output luminoso, ma hanno tempi di risposta relativamente bassi, mentre i secondi sono più veloci ma perdono in efficienza di luminescenza. Se a questo abbiniamo che l'alto Z dei costituenti e l'elevata densità degli scintillatori inorganici ben si prestano alla spettroscopia per raggi gamma, è ovvia la preferenza per questi ultimi nelle nostre applicazioni.

Il meccanismo di scintillazione nei materiali inorganici dipende dagli stati energetici determinati dal reticolo cristallino: come ci insegna la teoria delle bande, possiamo avere elettroni in banda di valenza legati ai siti del reticolo ed elettroni in banda di conduzione con sufficiente energia per migrare attraverso il cristallo.

L'assorbimento di energia permette all'elettrone di attraversare il band gap portandosi dalla banda di valenza a quella di conduzione lasciandosi dietro una lacuna, in seguito l'elettrone ritorna in banda di valenza rilasciando energia sottoforma di un fotone.

Purtroppo nel cristallo puro il meccanismo di rilascio di fotoni è inefficiente, inoltre il gap è così ampio che il processo rilascia fotoni con energia tale da non rientrare nello spettro del visibile.

Per aumentare la probabilità di emissione fotonica visibile sono aggiunte delle piccole quantità d'impurità (detti catalizzatori o attivatori), le quali introducono particolari siti reticolari, che si traducono in stati energetici intermedi all'interno della banda proibita nel quale l'elettrone eccitato si può collocare. Il minore gap energetico ora presente permette l'emissione di fotoni visibili e quindi il processo di scintillazione.

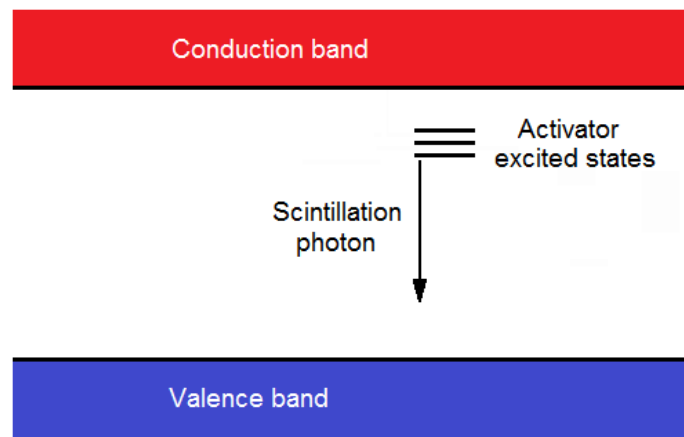


Figura 4.1: Schema a bande per un cristallo scintillatore attivato.

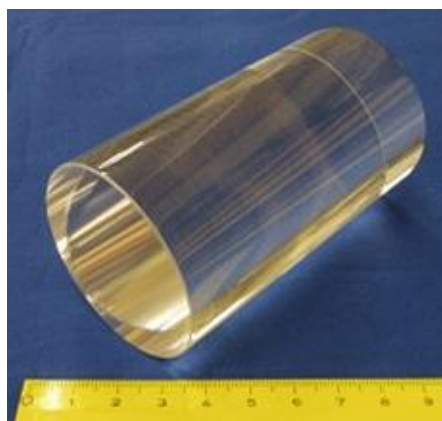
Un'altra importante conseguenza della luminescenza innescata con attivatori è che il cristallo diventa trasparente alla luce di scintillazione.

Nel cristallo puro invece è brutalmente liberata la stessa energia necessaria alla creazione della coppia elettrone-lacuna, con il risultato che gli spettri di emissione e assorbimento si sovrappongono dando luogo a un fenomeno di **autoassorbimento**.

Lo spettro di emissione di uno scintillatore inorganico è caratterizzato da un picco in corrispondenza di una determinata lunghezza d'onda, bisogna quindi aver cura che i dispositivi utilizzati per la rilevazione di luce abbiano la loro zona di massima sensibilità quanto più vicina possibile a questo picco.

Tra gli scintillatori inorganici più utilizzati ricordiamo gli alogenuri alcalini (tra i quali vi è il comunissimo NaCl), composti ionici costituiti da un metallo alcalino e un alogeno, ma anche il BGO (Bismuto Germanato) e il LYSO.

Ed è proprio il LYSO (Cerium-doped Lutetium-Yttrium Orthosilicate) il materiale utilizzato come scintillatore nel nostro caso; riportiamo di seguito alcune sue caratteristiche.



Picco di emissione (nm)	420
Numero atomico effettivo	65
Indice di rifrazione	1.82
Light Yield (fotoni/MeV)	30000

Figura 4.2: Cristallo di LYSO e tabella con le specifiche ottiche.

Più che l'autoassorbimento (il quale è la principale causa di perdite nella raccolta di luce solo per cristalli molto grandi), i meccanismi che principalmente intervengono nella diminuzione di efficienza

della raccolta di luce emessa sono le perdite all'interfaccia dello scintillatore, quindi l'uniformità nella raccolta della luce dipende principalmente dalle condizioni al contorno tra lo scintillatore e l'ambiente in cui è immerso.

La luce di scintillazione è emessa in ogni direzione e solamente una limitata frazione di essa viaggia direttamente verso il fotorivelatore, i restanti fotoni, prima di essere raccolti, subiranno una o più riflessioni sulla superficie del cristallo.

Fatta questa premessa, possono presentarsi due situazioni quando un fotone raggiunge l'interfaccia:

- Se l'angolo d'incidenza è maggiore dell'angolo critico θ_c si avrà riflessione totale (caso 1 in figura).
- Se l'angolo d'incidenza è minore di θ_c , avremo riflessione parziale (detta riflessione di Fresnel) e trasmissione parziale (caso 2 in figura).

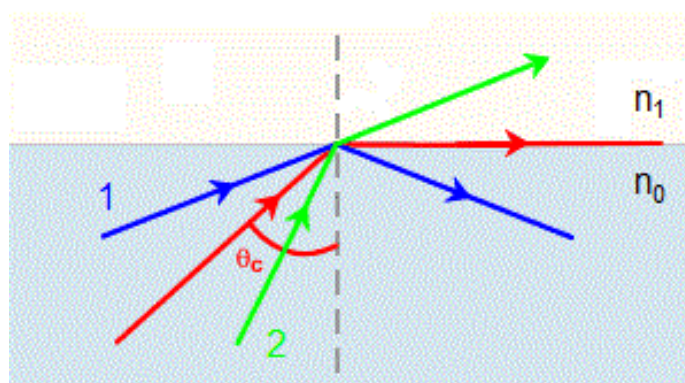


Figura 4.3: Comportamento all'interfaccia dello scintillatore.

L'angolo critico è determinato dagli indici di rifrazione del cristallo e del mezzo circostante secondo la formula:

$$\theta_c = \sin^{-1} \frac{n_1}{n_0}$$

Per catturare anche la luce che normalmente andrebbe persa alle superfici, l'accorgimento che adottiamo consiste nel rivestire il cristallo di un materiale riflettente (nel caso particolare Teflon).

Lo scopo ultimo è quello di raccogliere la maggior parte di luce emessa, questo per due motivi:

- Una riduzione del numero di fotoni contribuenti all'impulso misurato comporta un allargamento dello spettro energetico di risposta, in altre parole un peggioramento nella risoluzione energetica.
- Una raccolta di luce perfettamente uniforme assicura che tutti gli eventi depositanti la stessa energia, diano luogo a impulsi della stessa ampiezza, indipendentemente dal punto del cristallo in cui scaturisce l'evento.

Capitolo V:

Silicon PhotoMultiplier

Un fotodiodo è un particolare tipo di giunzione $p-n$ con drogaggio molto asimmetrico, cioè tipicamente la zona p è molto più drogata di quella n .

Tali dispositivi se polarizzati direttamente non si distinguono in alcun modo da un diodo normale, è interessante invece la polarizzazione inversa, la quale provoca la formazione di una regione di svuotamento, detta *depletion region*, dalle dimensioni lineari date dalla formula:

$$W = \sqrt{2 \epsilon \frac{V + V_{bi}}{N e}} = \sqrt{2 \rho \mu \epsilon (V + V_{bi})}$$

dove V è la polarizzazione applicata esternamente, V_{bi} è la cosiddetta tensione *built in*, creata internamente, N è la concentrazione di drogaggio, e la carica elettrica, ϵ la costante dielettrica del materiale, ρ è la resistività e μ la mobilità dei portatori di carica.

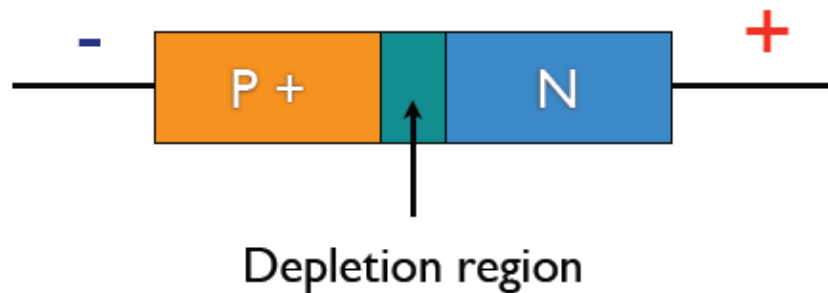


Figura 5.1: Schema semplificato di un fotodiode.

Se l'energia $h\nu$ del fotone incidente nella *depletion region* è maggiore del salto tra banda di valenza e banda di conduzione del materiale, avviene la creazione di una coppia elettrone-lacuna ($e-h$ in breve) e, a causa del campo elettrico presente, l'elettrone si sposterà verso la regione n e la lacuna verso la zona p .

La mancanza della coppia nella regione di svuotamento produce una fotocorrente inversa che è letta come segnale elettrico.

È molto importante la geometria dello strumento, che deve essere progettata per massimizzare la probabilità d'incidenza del fotone nella regione di svuotamento.

Le caratteristiche tecniche principali di un fotodiode in termini di risposta ed efficienza sono la *quantum efficiency* e la *responsivity*: la prima si definisce come il numero di coppie $e-h$ generate per ciascun fotone incidente, mentre la seconda rappresenta il rapporto tra la fotocorrente prodotta e la potenza ottica incidente.

Il guadagno interno dei fotodiodi in polarizzazione inversa è pressoché nullo.

Per le applicazioni nella Fisica delle Alte Energie è più comune l'utilizzo dei fotodiodi a valanga (*Avalanche PhotoDiode*, APD), che si differenziano per alcuni aspetti dai precedenti.

Un fotodiodo a valanga può essere costituito da quattro strati di semiconduttore drogati in maniera differente:

- Una zona $p+$ con grande concentrazione di accettori (accettori/ $cm^3 > 10^{17}$).
- Una zona di semiconduttore intrinseco, cioè non drogato, che serve a mantenere quasi costante il campo elettrico all'interno della giunzione e a diminuirne la capacità.
- Una zona p drogata meno intensamente della precedente.
- Una zona $n+$, dove la concentrazione di donatori è grande (numero di donatori/ $cm^3 > 10^{17}$).



Figura 5.2: Schema semplificato di un APD.

Le cariche primarie prodotte attraversano la zona drogata p e producono delle cariche secondarie, che daranno luogo alla fotocorrente, cruciale è dunque l'inserimento di tale strato, poiché la generazione a valanga è dovuta alla sua presenza.

I parametri indicativi del fotodiodo a valanga sono il *random multiplier factor* o *guadagno* e il *noise factor*: il primo, indicato con la lettera M , rappresenta il numero di coppie secondarie $e-h$ generate per ciascuna coppia primaria, mentre il fattore di rumore, indicato con F è definito come

$$F = \frac{\langle M^2 \rangle}{\langle M \rangle^2}$$

Il guadagno di un APD cresce con la polarizzazione inversa e raggiunge tipicamente il valore di $10^2 \div 10^3$, per raggiungere un regime di maggiori guadagni è necessario operare a tensioni maggiori del cosiddetto *break down voltage* (tensione di rottura), oltre il quale si vede nella caratteristica del diodo un incremento esponenziale della corrente.

Per evitare danneggiamenti dello strumento sono utilizzate tecniche di contenimento della valanga dette di *quenching*, in modo tale da poter aumentare la tensione.

Un APD che lavora nel modo appena descritto è detto operante in *Geiger mode* e può raggiungere guadagni compresi fra 10^4 e 10^7 .

Una matrice di APD operanti in regime *Geiger* costituisce un Silicon PhotoMultiplier (SiPM); tutti questi APD hanno in serie una resistenza (*quenching* passivo) e sono collegati in parallelo tra di loro nello stesso substrato di Silicio.

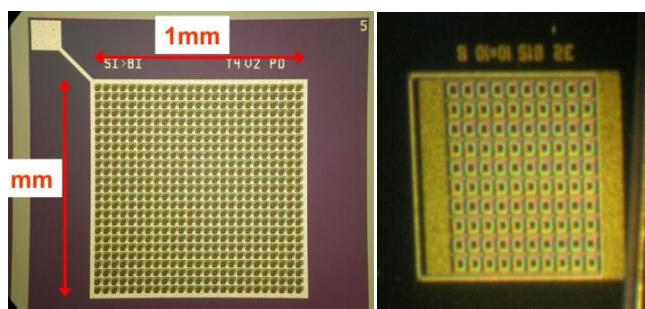


Figura 5.3: Ingrandimenti al microscopio di due SiPM.

La tensione operativa è del $10 \div 20\%$ superiore al valore di *breakdown* e non può essere incrementata a piacimento poiché il rumore di fondo è proporzionale alla tensione di polarizzazione inversa delle giunzioni; i valori tipici si aggirano intorno ai $30 \div 70$ V.

Il segnale in uscita da un SiPM è la somma analogica dei segnali degli APD; in questo senso la matrice può essere considerata uno strumento analogico: mentre un APD è un dispositivo binario, il SiPM fornisce un segnale elettrico proporzionale al numero di fotoni incidenti se questi sono in numero minore delle microcelle di APD (altrimenti si raggiunge la saturazione), che chiameremo pixel. In questo senso è molto importante valutare le dimensioni dei pixel e l'area morta che li circonda secondo l'uso che si vuole fare dello strumento e dell'illuminazione prevista.

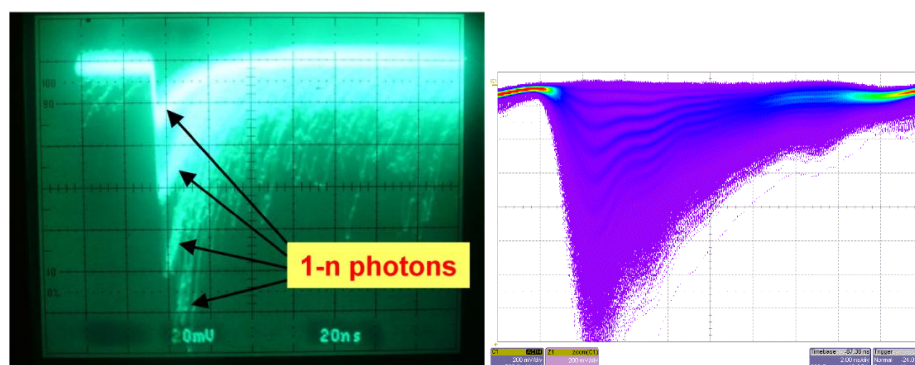


Figura 5.4: Segnale del SiPM all'oscilloscopio (da sinistra) analogico e digitale.

Il valore tipico per il tempo di salita, dovuto alla scarica Geiger, si aggira intorno al ns , mentre il tempo di discesa assume valori tipici di alcune decine di ns .

Da questi valori si evince che lo strumento è caratterizzato da un basso valore di tempo di recupero, il che lo rende appropriato per applicazioni con alte frequenze di conteggi.

Il guadagno tipico di un fotomoltiplicatore al Silicio è comparabile a quello dei PM tradizionali e assume un valore di circa $10^5 \div 10^6$, tuttavia è stata riscontrata una dipendenza dalla temperatura e dalla tensione di polarizzazione inversa applicata che segue approssimativamente il seguente andamento

$$\frac{dG}{G} \approx 7 \frac{dV}{V} \quad \frac{dG}{G} \approx 1.3 \frac{dT}{T}$$

Il rumore di questo strumento è principalmente dovuto alla cosiddetta *dark current* (corrente oscura), provocata dalla creazione termica di portatori di carica nella *depletion region* del fotodiodo, che innescano la scarica Geiger esattamente come un fotone

incidente e non possono pertanto essere distinti da un segnale reale. È molto importante, dunque, conoscere l'andamento specifico del rumore per poterlo minimizzare mediante la scelta dei parametri di lavoro.

Il rumore è proporzionale alla tensione e alla temperatura e, poiché il guadagno ha un andamento simile, è necessario bilanciare i valori cui lavorare per ottenere le condizioni ottimali.

L'efficienza dei SiPM è ridotta dallo spazio morto tra i pixel che compongono la matrice: un fotone incidente in questa zona, infatti, è perso e non provoca alcuna reazione nello strumento.

Il parametro fondamentale in questo senso è il *fill factor*, che è dato dal rapporto tra l'area attiva di un pixel e la sua area totale compresi gli elementi circuitali: maggiore è il fattore di riempimento, maggiore sarà l'efficienza dello strumento.

Importante in questo senso è anche il cosiddetto fenomeno di *crosstalk* tra le celle della matrice; durante la scarica Geiger, infatti, sono prodotti dei fotoni che possono migrare verso una cella adiacente e provocare in essa una scarica spuria.

I rimedi possibili sono due: la riduzione del guadagno oppure l'isolamento ottico dei singoli pixel; in questa seconda scelta è tuttavia necessario evitare la riduzione del *fill factor* per non intaccare ulteriormente l'efficienza dello strumento.

Da un confronto fra i SiPM e i tubi fotomoltiplicatori (PMT) emergono le seguenti considerazioni:

- La tensione di lavoro dei SiPM è un aspetto molto favorevole: tensioni di alcune decine di Volt si contrappongono, infatti, ai valori nominali di circa 800 V per i PMT.
- Le dimensioni delle matrici di fotodiodi sono molto minori, si aggirano intorno al mm^2 e li rendono pertanto più maneggevoli.
- Al contrario dei PMT tradizionali, i SiPM non sono disturbati dai campi magnetici.

Capitolo VI:

Strumentazione

Tutta la strumentazione utilizzata per eseguire le misure può essere schematizzata in quattro macroblocchi funzionali:

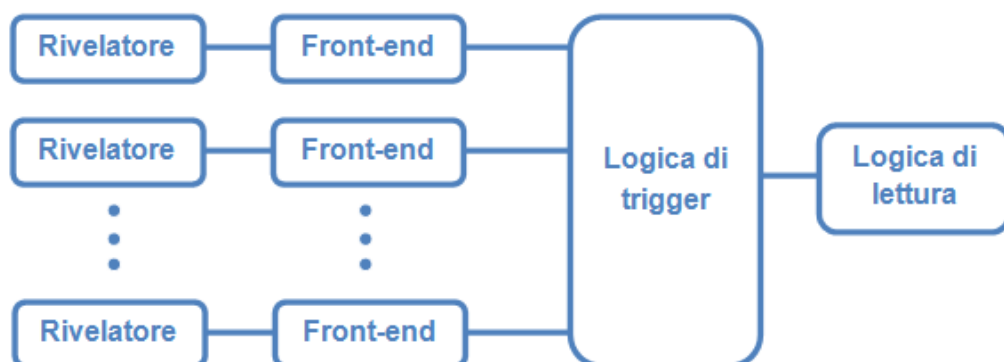


Figura 6.1: Schema a blocchi generico per le misure effettuate.

- *Rivelatore*: ha il compito di tradurre la rivelazione di particelle in un segnale, la cui natura dipende dalle informazioni che si vogliono ottenere e dal tipo di energia delle particelle da rivelare.

- *Elettronica di front-end*: serve a trattare i segnali uscenti dal rivelatore, così da ricavarne le informazioni desiderate. Con il termine trattare intendiamo l'effettuare sul segnale operazioni utili per la lettura dei dati, quali amplificazione, discriminazione, shaping e conversione.
- *Logica di trigger*: seleziona quali delle molteplici informazioni che arrivano dai vari front-end sono veramente utili per ottenere le informazioni desiderate.
- *Logica di lettura*: permette di inviare i dati che arrivano dalla logica di trigger ai successivi stati di acquisizione, rendendoli comprensibili (formattazione) e compatibili (temporizzazione.)

Scendendo più nel dettaglio, al di là del rivelatore che nel nostro caso è sempre un fotorivelatore, tutta la strumentazione elettronica utilizzata negli esperimenti di fisica nucleare segue quasi sempre gli standard **NIM** (Nuclear Instrument Module) e **CAMAC** (Computer Automated Measurement and Control).

Lo standard NIM è stato il primo e tutt'oggi il più semplice standard introdotto nella fisica nucleare, esso, oltre a definire delle specifiche meccaniche (dimensioni dei moduli, tipi di cavi usati, ecc) ed elettriche (forma e ampiezza del segnale, tensioni di alimentazione, ecc), introduce il concetto di *sistemi elettronici a moduli*, offrendo enormi vantaggi in termini di flessibilità,

intercambiabilità della strumentazione, facilità di aggiornamento e manutenzione.

In pratica un sistema elettronico NIM è costituito da moduli alloggiati in una struttura, detta *crate*, la quale provvede alla loro alimentazione con tensioni continue differenziali da $\pm 24\text{V}$, $\pm 12\text{V}$ e $\pm 6\text{V}$.

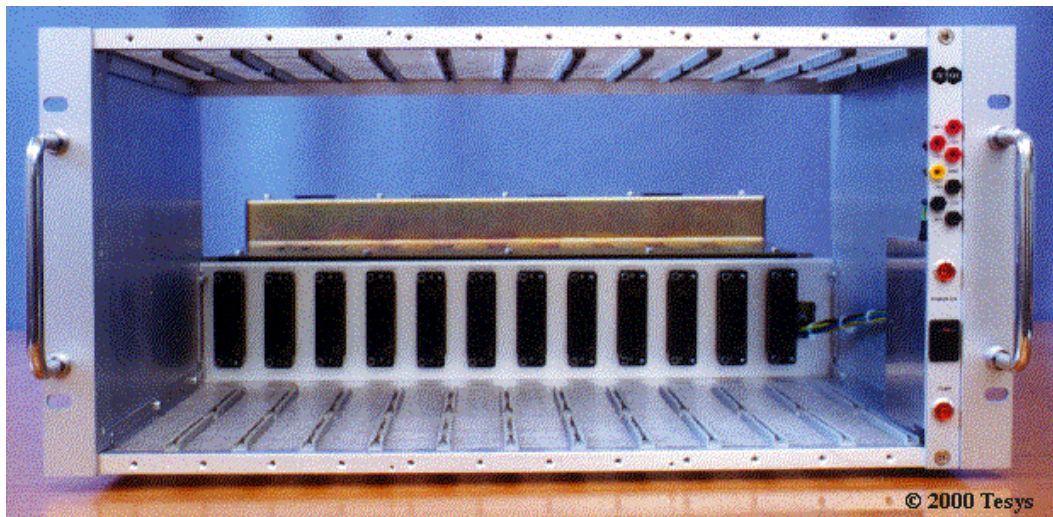


Figura 6.2: Crate NIM.

Ogni crate può alloggiare fino a 12 moduli di larghezza singola (esistono anche moduli di larghezza doppia e tripla), ciascuno dei quali ha una funzione specifica (amplificatori, ADC, DAC, ecc).

Per quanto riguarda il supporto meccanico lo standard NIM prevede l'utilizzo di cavi schermati RG58 ed RG174 provvisti di connettori LEMO.

Inoltre è introdotto anche lo standard logico NIM, una logica in corrente con uno stato “true” negativo (-16 mA su $50\ \Omega$, cioè -0.8 V) e uno stato “false” nullo.

Il CAMAC invece è uno standard bus per l’acquisizione di dati ed è stato introdotto come estensione dello standard NIM per permettere il controllo computerizzato della strumentazione; anch’esso è un sistema modulare con 25 slot di alloggiamento.

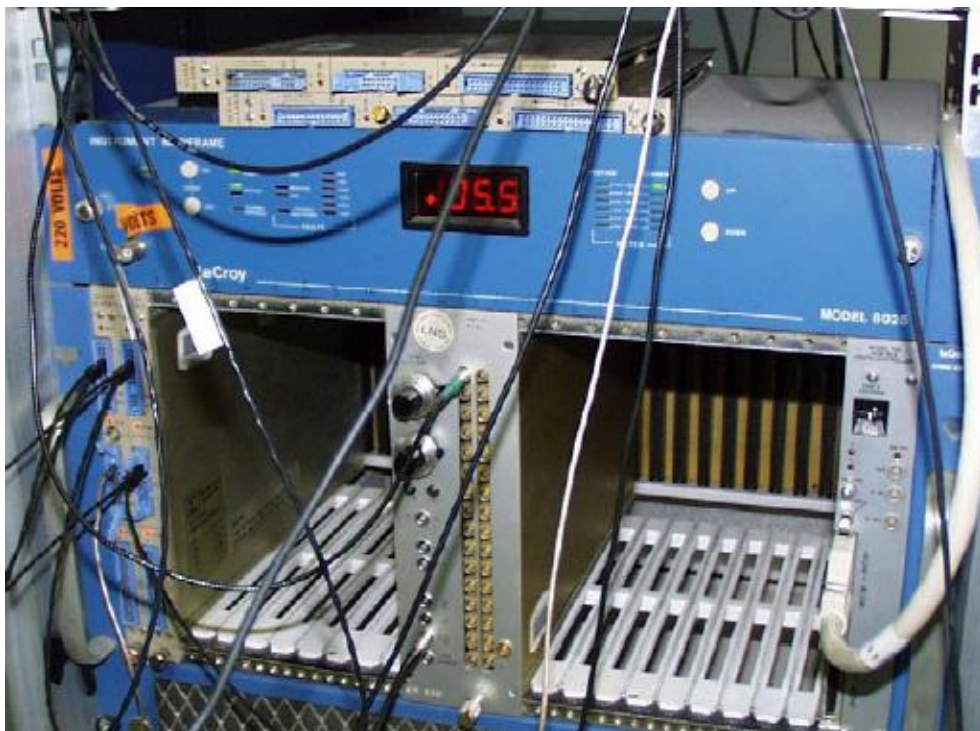


Figura 6.3: Crate CAMAC.

Il CAMAC si differenzia dal NIM nel contenere non solo i segnali di alimentazione, ma anche un bus di dati gestito dal *crate controller* (CC), il quale funge da pilota per il controllo dei

comandi e dei dati verso le altre schede del crate, nonché da interfaccia verso il computer.

In particolare in un CAMAC bus esistono tre categorie di segnali:

- *Alimentazioni*: banalmente alimentano i moduli contenuti nel crate con gli stessi valori di tensione visti nel NIM.
- *Segnali comuni*: dati, indirizzi e segnali di controllo necessari per il funzionamento del sistema.
- *Segnali punto-punto*: linee non condivise tra le varie schede, ma dedicate, che da un singolo slot vanno al CC.

A quest'ultima categoria appartengono il *crate address* (N) e il *look at me* (LAM).

Il segnale N memorizza gli slot in cui sono inseriti i moduli, cosicché, quando il CC riceve un segnale di controllo, seleziona correttamente una determinata posizione in memoria nel crate address, prelevando i dati contenuti nel rispettivo modulo.

La linea di LAM è invece gestita generalmente da un segnale esterno al CAMAC che ha lo scopo di fornire al controller una richiesta di attenzione.

Questo compito è svolto nel caso particolare dal modulo **TINA** (Integrated Trigger for Nuclear Acquisition data), il quale, appunto, invia segnali di attenzione al controller per registrare i dati contenuti nei convertitori.

Il TINA genera inoltre un segnale in logica NIM, chiamato INHIBIT, che ha la funzione di mantenere bloccato il sistema di

acquisizione durante la fase di elaborazione e lettura dati e che quindi viene gestito dal TINA stesso secondo un preciso schema temporale.

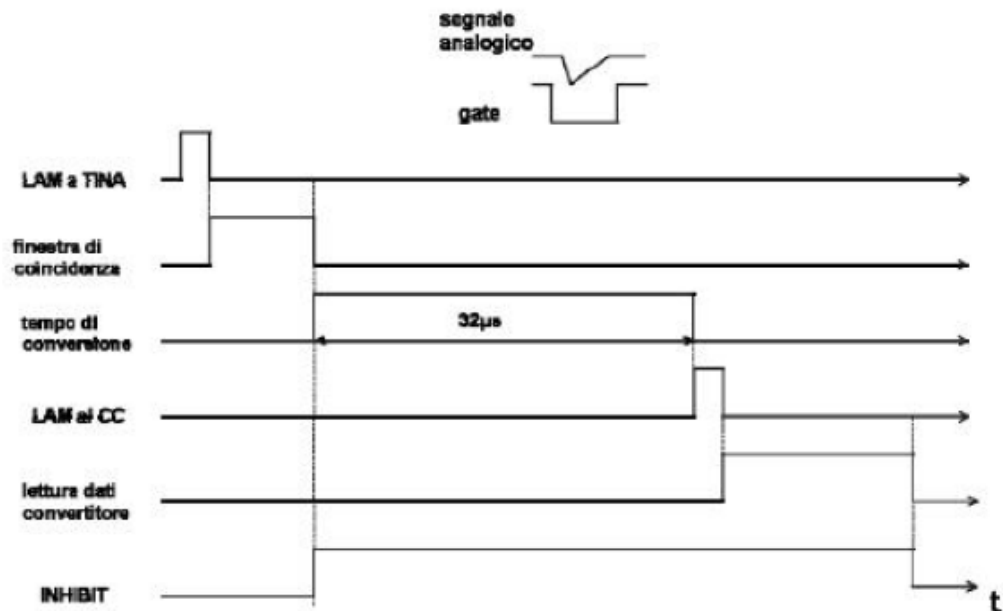


Figura 6.4: Temporizzazione dei segnali logici tra TINA e CAMAC.

Per regolare i segnali di LAM e INHIBIT occorre costruire una logica di trigger tra i segnali che viaggiano nei moduli NIM e poi sincronizzarli con il segnale LAM inviato al CC, cosicché, quando la logica di trigger decide che nei moduli posizionati nel CC vi sono dati utili, il segnale LAM viene inviato automaticamente per prelevarli.

Il LAM gioca dunque un ruolo fondamentale nell'interfacciare l'hardware e il software necessari ai fini delle misure.

Parlando di software bisogna aggiungere che tutte le analisi di post-processing sui dati raccolti sono state svolte con il **PAW** (Physics

Analysis Workstation), un software sviluppato al CERN basato su un'estesa collezione di librerie Fortran.

Capitolo VII:

Misure di risoluzione temporale

Le misure effettuate sono mirate alla caratterizzazione di SiPM di tre diversi produttori: Hamamatsu, SensL e ST Microelectronics, innanzitutto dal punto di vista della risoluzione temporale del singolo dispositivo, al fine di valutare l'idoneità per l'impiego nelle misure con gli scintillatori di eventi in coincidenza in presenza di una sorgente radioattiva.

Bisogna aggiungere però che il sensore di STM è stato testato solamente per completezza poiché non è attualmente prevista la produzione su grandi numeri.

Riportiamo di seguito le caratteristiche dei dispositivi presi in esame.

	Hamamatsu	SensL	STM
Area dispositivo	9 mm ²	9 mm ²	0.25 mm ²
Numero di celle	14400	3640	100

Dimensioni cella	25 x 25 μm^2	35 x 35 μm^2	50 x 50 μm^2
V breakdown	70 V	28 V	30 V
Picco spettrale	440 nm	490 nm	-

N.B. Non vengono riportati i valori di guadagno e dark count in quanto aumentano con la tensione di polarizzazione.

Esaminiamo ora il sistema di misura utilizzato per l'acquisizione degli spettri di carica e tempo, scendendo nel dettaglio sulle caratteristiche dei moduli scelti e sulle motivazioni che stanno dietro la preferenza di un modulo rispetto ad un altro.

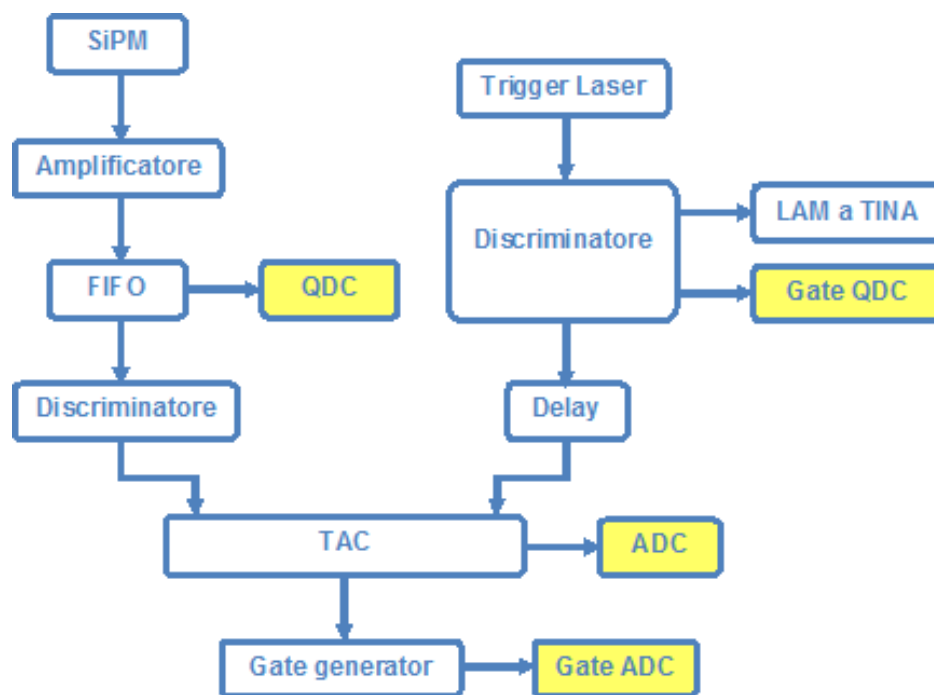


Figura 7.1: Schema elettronico usato per l'acquisizione degli spettri.

- *Amplificatore Ortec FTA410*

L'amplificazione è necessaria per portare il segnale da valori intorno al centinaio di μV fino a qualche decina di mV ; in particolare questo amplificatore è caratterizzato da un guadagno di circa 200, un'ottima risposta temporale con un tempo di salita inferiore al nanosecondo (il che è ottimo per applicazioni come questa, in cui abbiamo segnali molto veloci) e un rumore di fondo relativamente basso (inferiore a $20 \mu\text{V rms}$).

- *Discriminatore LeCroy 4608C*

Sulla scelta del discriminatore va aperta una parentesi un po' più ampia, in quanto i *fast-timing discriminator*, ossia quelli che ci permettono di contare elevati tassi di rapidi impulsi identificandone con precisione il tempo di arrivo, si suddividono in due categorie.

I **LED** (che sta per *leading-edge discriminator*) sono costituiti da un semplice comparatore di tensione riferito ad una soglia variabile dall'utente, si ha così un'indicazione precisa del tempo di arrivo di ogni impulso.

Questo è vero fintanto che non interviene il rumore a causare un'incertezza, detta *jitter*, dovuta all'attraversamento rumoroso del valore di soglia; il contributo del rumore al *timing jitter* segue una legge del tipo

$$\frac{e_n}{dV/dt}$$

Al numeratore abbiamo l'ampiezza del rumore e al denominatore la pendenza del segnale, quindi se vogliamo minimizzarlo basta impostare la soglia nel punto di massima pendenza.

Oltre al jitter vi è un altro fenomeno che inficia la risoluzione temporale, parliamo di una dipendenza dall'ampiezza del segnale in ingresso detta *walk*.

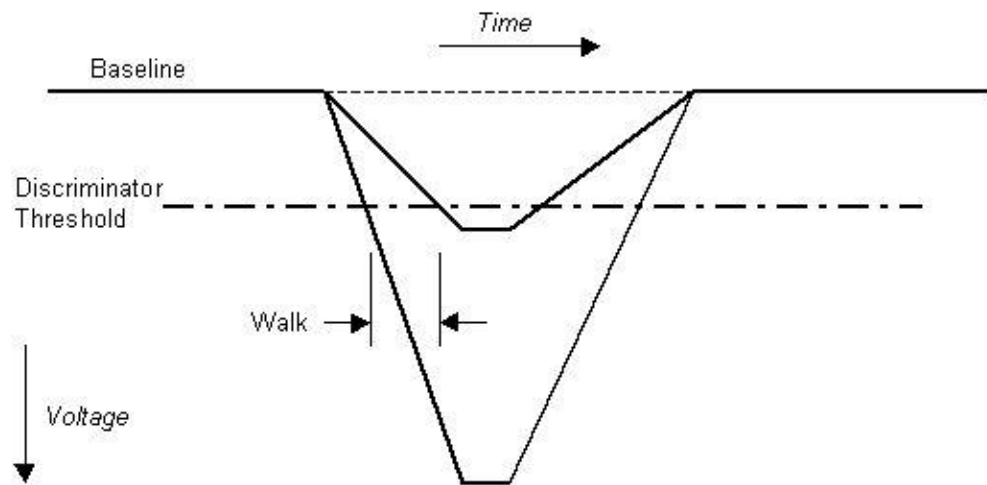


Figura 7.2: Causa della presenza di walk.

Come si può osservare in figura, la pendenza finita del segnale comporta che due segnali di stessa forma ma diversa ampiezza attraversino la soglia in due istanti differenti: quello più piccolo prima e quello più ampio dopo.

La differenza tra gli istanti di attraversamento è detta appunto walk. Questo degrada la risoluzione temporale, specialmente quando si processano, come nel caso dei segnali in uscita dai SiPM, molti tipi diversi di impulsi.

Si osserva però che esiste un intervallo di ampiezze dei segnali in ingresso al LED per le quali la risoluzione temporale risulta ottimizzata, questo spinge verso la realizzazione di un circuito che faccia il trigger automatico della soglia migliore: nasce così il **CFD** (*constant fraction discriminator*).

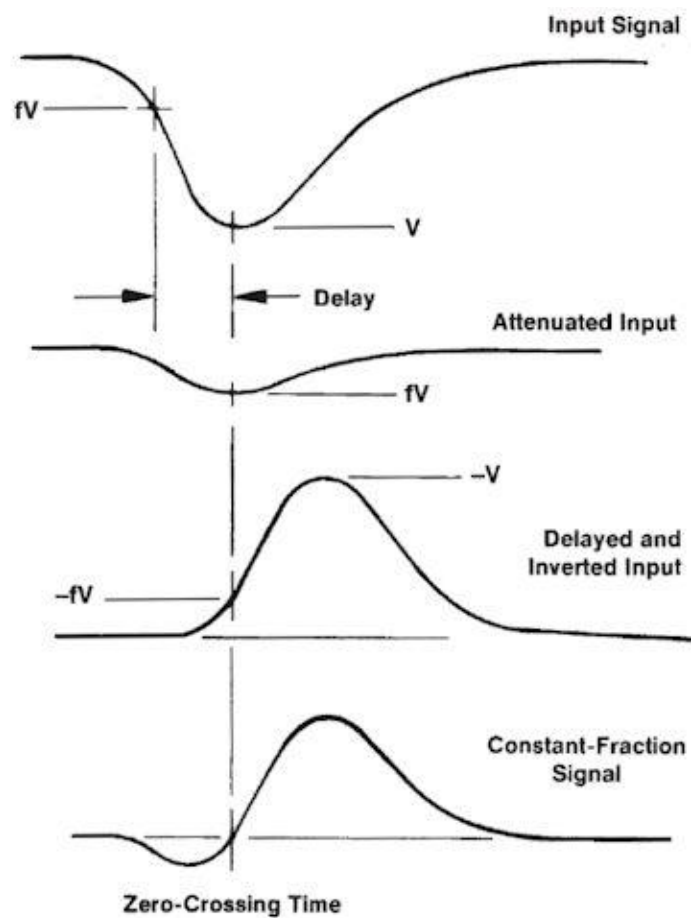


Figura 7.3: Realizzazione del segnale Constant-Fraction.

Banalmente il circuito può essere schematizzato in due rami: uno attenua il segnale in ingresso, l'altro lo inverte e lo ritarda; sommando i segnali sui due rami otteniamo così il constant fraction signal.

Il CFD presenta però un maggior jitter del LED, per questo motivo, e dato anche che la correzione del walk è una banale pratica di post processing, abbiamo preferito per la nostra applicazione il LED.

La soglia del discriminatore è stata scelta sufficientemente alta per eliminare il rumore e i conteggi scorrelati di eventi “di zero”, ossia quegli eventi in cui, lavorando a bassa intensità, il sensore non rivela fotoni e la carica che arriva al convertitore è solo quella del livello di riferimento, ma sufficientemente bassa per non perdere eventi correlati con conseguente peggioramento del timing.

- *Linear FIFO (fan in/fan out) LeCroy 428F*

Quello che nello schema è indicato con FIFO altro non è che uno “sdoppiatore”, ossia un modulo che, dato un segnale in ingresso, ci permette di prelevarlo su più uscite contemporaneamente e senza attenuazioni, ma questo a prezzo di un piccolo aumento nel rumore. Inoltre il modulo ci permette di modificare lo *zero level* del segnale, cosa che risulta molto utile in quando il LED ha una soglia minima di 25 mV, ma per i segnali in gioco abbiamo bisogno di una soglia inferiore.

- *Unità laser ALS PILAS EIG1000D*

L’unità di controllo del laser ha due scopi nel setup sperimentale proposto, il primo è quello di illuminare il SiPM con impulsi laser blu a 408 nm (per emulare quanto più possibile il comportamento dello scintillatore, che ha un picco di emissione di a 420 nm) di durata 50 ps, e il secondo è quello di fornire un segnale di trigger in

logica ECL, che poi andrà convertito in NIM, il quale funge sia da LAM per il TINA che da segnale di “start” per il TAC (questo verrà discusso più avanti).

- *Time to Pulse Height Converter Ortec 457*

Il TAC è il punto nevralgico delle misure di risoluzione temporale: al suo ingresso inviamo un segnale di start proveniente dal laser e un segnale di stop proveniente dal SiPM, opportunamente ritardati con dei moduli passivi di delay formati da cavi avvolti (circa 1 ns ogni 5 cm) per far sì che la dinamica di uscita del TAC non sia né troppo piccola né satura.

Questo ci dà fisicamente una misura del tempo che intercorre tra l’illuminazione del SiPM e l’arrivo del segnale in uscita, sebbene non siamo interessati a questo intervallo di tempo in sé, ma alla distribuzione statistica degli intervalli di tempo.

- *Silena ADC 4418/Q*

Tale modulo, alloggiato nel crate CAMAC, si occupa della conversione del segnale in carica del SiPM in uno spettro digitale.

Per far ciò esso deve ricevere il segnale analogico su uno dei suoi 10 canali, nonché un segnale di *gate* in logica ECL tale da contenere il segnale, nel caso particolare il gate è generato dal canale del discriminatore che riceve il trigger del laser e viene regolato su una durata dell’ordine di qualche decina di ns (non è necessaria una elevata precisione, basta che il segnale di gate contenga, in termini di durata, il segnale da convertire).

- *Silena ADC 4418/V*

Il funzionamento è analogo al QDC, con la differenza che questo modulo riceve il segnale del TAC (diventa a tutti gli effetti un TDC) in uno dei suoi canali, inoltre il gate, dato che i segnali uscenti dal TAC hanno durate ben più elevate di quelle dei SiPM, viene generato da un multivibratore astabile (*Dual Gate Generator LeCroy 222*) che collega in serie due generatori di impulso di gate per averne uno solo regolabile sia in salita che in discesa.

Si ottiene così un gate della durata di qualche centinaio di ns (anche qui la precisione non è di primaria importanza).

Il sistema così collegato deve seguire una breve procedura di calibrazione, la quale permette di ricavare sia il contributo del solo sistema elettronico all'indeterminazione temporale, sia quanti ps corrispondono ad una unità (detta anche canale) nell'istogramma del software di acquisizione.

La procedura consiste nel sostituire in prima battuta il segnale del SiPM con il trigger del laser, in modo da prelevare, con incrementi di delay ben definiti, uno spettro “a pettine”, ossia uno spettro costituito da picchi distanziati tra di loro di una quantità temporale nota, ciascuno con la minima indeterminazione temporale.

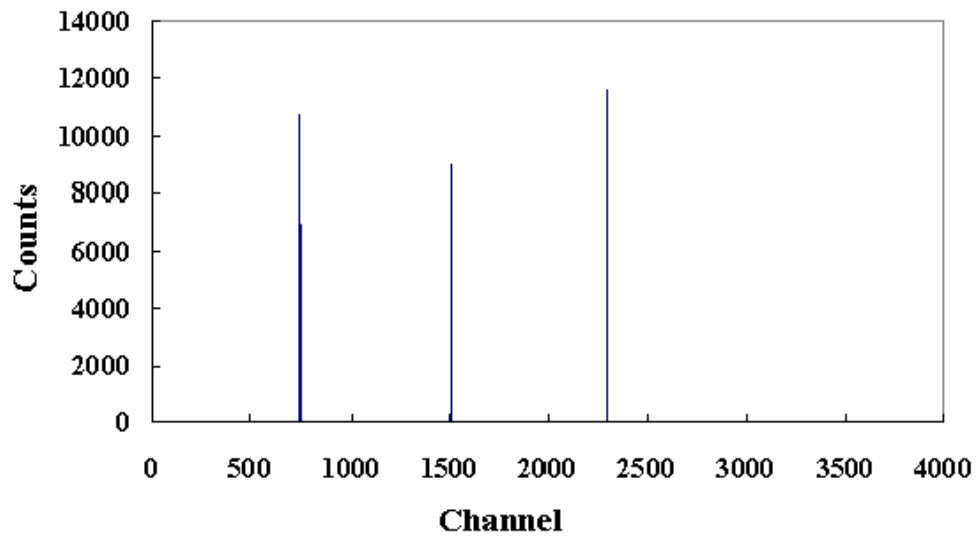


Figura 7.4: Spettro per la calibrazione del TAC.

Interpolando i punti si ottiene una retta il cui inverso del coefficiente angolare ci da:

$$\# \left(\frac{psec}{canale} \right) \sim 21.5$$

Avendo questo risultato possiamo andare a valutare la FWHM dei picchi ottenendo una risoluzione per il ramo del trigger del laser di 34 psec.

Dato che però abbiamo due rami costituiti da elettronica diversa, avremo due indeterminazioni temporali diverse che andranno in somma quadratica.

Da un'acquisizione fatta sostituendo il trigger del laser con il segnale del SiPM si ottiene infatti una FWHM di 134 psec, da cui si ricava il contributo dell'elettronica all'indeterminazione:

$$\sqrt{\frac{(34)^2}{2} + \frac{(134)^2}{2}} \cong 100 \text{ psec}$$

Tale valore rappresenta un estremo inferiore per la risoluzione temporale, che potrà essere solamente superiore ad esso.

Precisiamo però che questi valori sono privi di correzione del walk sui dati (procedura che discuteremo più avanti) e quindi in realtà i valori di incertezza sono ancora più bassi.

Bisogna approfondire adesso quello che nello schema è indicato semplicemente con il blocco SiPM, esso è costituito sì dal fotosensore e dalla sua alimentazione, ma anche dal seguente circuito di readout.

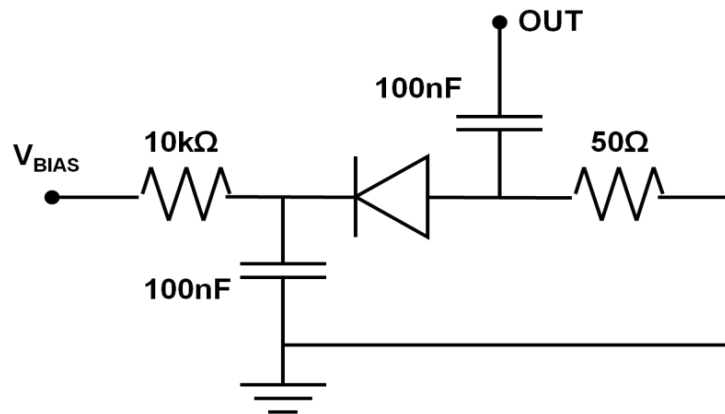


Figura 7.5: Circuito di polarizzazione dei SiPM.

Durante la fase statica tutta la tensione V_{BIAS} cade ai capi del SiPM, invece nella fase dinamica le capacità si cortocircuitano costituendo al catodo un percorso preferenziale verso massa e all'anodo un

passaggio per il segnale che cade ai capi della resistenza da $50\ \Omega$ (valore scelto per garantire il massimo trasferimento di potenza).

Inoltre si è provveduto a due requisiti ambientali di primaria importanza: isolamento elettromagnetico e isolamento da luce ambientale, in quanto in primo luogo il fotosensore non può ricevere un'illuminazione diversa da quella che gestiamo noi con il laser, pena la compromissione delle misure, e secondariamente si osservano all'oscilloscopio fenomeni di pesante interferenza elettromagnetica.

Si è realizzata allora una rudimentale ma alquanto efficace gabbia di Faraday, l'intero sistema è stato coperto da un telo nero e le misure sono state effettuate al buio.



Figura 7.6: Setup sperimentale per la schermatura elettromagnetica (sinistra) e ottica (destra).

Il sistema è dunque pronto per procedere con l'acquisizione degli spettri di carica e tempo in condizioni di:

- Bassa intensità del laser (il nostro scopo è visualizzare i picchi dei singoli fotoni).
- Frequenza di impulsi 500 Hz.
- Overvoltage pari al 2% della V_{br} nominale.

Parlando dei discriminatori, ma anche quando abbiamo determinato il contributo dell'elettronica al rumore, si è fatto cenno alla procedura di correzione del walk, causato dal LED, sui dati sperimentali ottenuti.

Per procedere alle correzioni bisogna dapprima realizzare un istogramma di dispersione carica-tempo per ciascuno dei sensori.

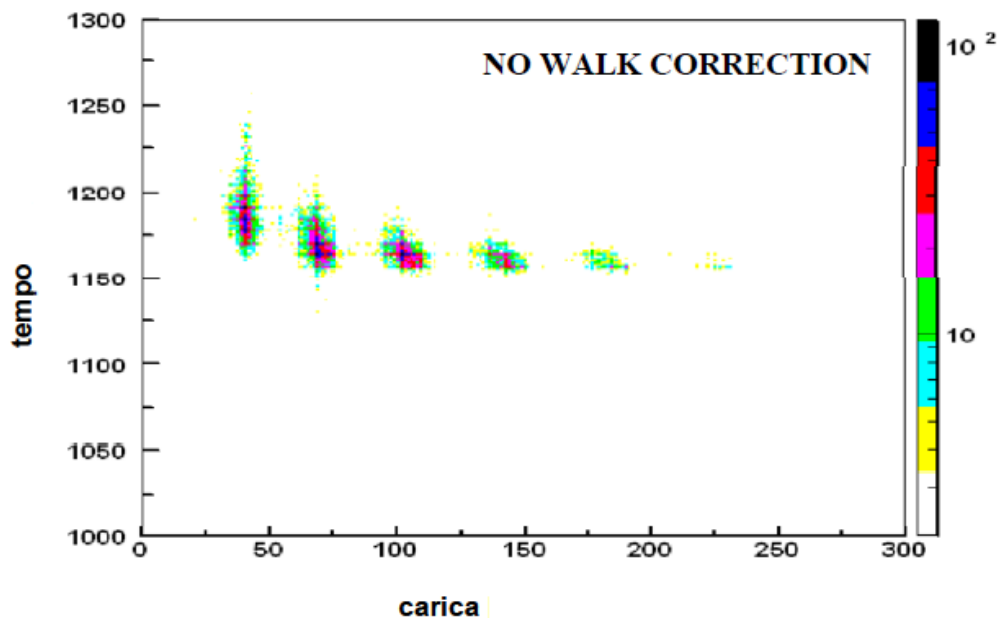


Figura 7.7: Scatter plot tempo-carica senza correzione del walk.

Come previsto il walk è più evidente per segnali di ampiezza minore, infatti le isole sono spostate verso tempi più lunghi a causa dell'ampiezza del segnale, troppo vicina alla soglia del discriminatore.

Per riuscire a valutare la giusta correzione da apportare ai dati si è seguita una procedura particolare: dato che nel grafico bidimensionale i punti sperimentali si mostrano come delle isole, possiamo distinguere i conteggi di tempo in base alla carica rivelata; a questo punto basta riportare i centroidi tempo in funzione della carica, espressa come numero di fotoni al quale fa riferimento il centroide.

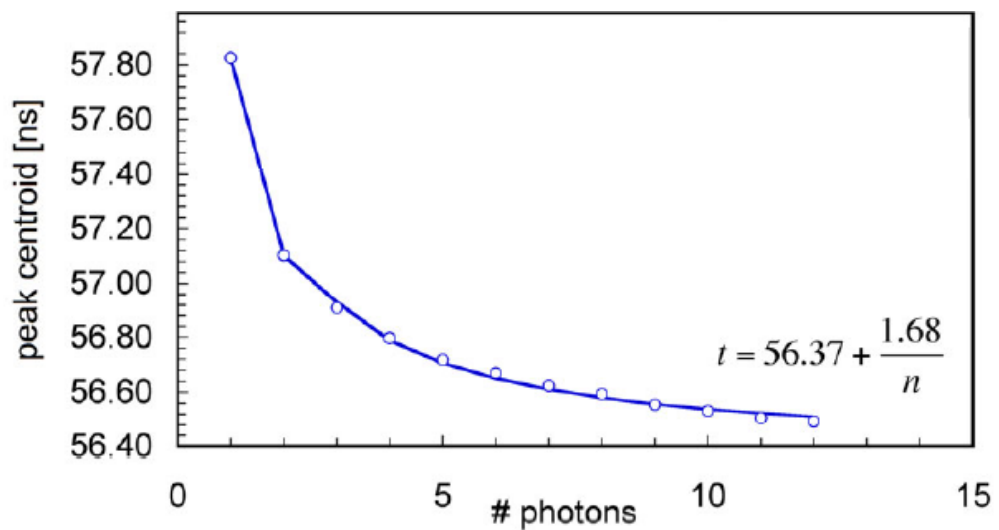


Figura 7.8: Interpolazione per la correzione del walk.

La curva risultante è un ramo di iperbole, il che è perfettamente in accordo con quello che ci aspettavamo, dato che la pendenza del

segnale in salita è $m \cdot n$, dove m è il coefficiente angolare del segnale a singolo fotone e n è il numero di fotoni, quindi l'istante temporale di attraversamento della soglia segue una legge inversa:

$$Y = \frac{a}{X} + b$$

che è proprio la formula da applicare per correggere i dati.

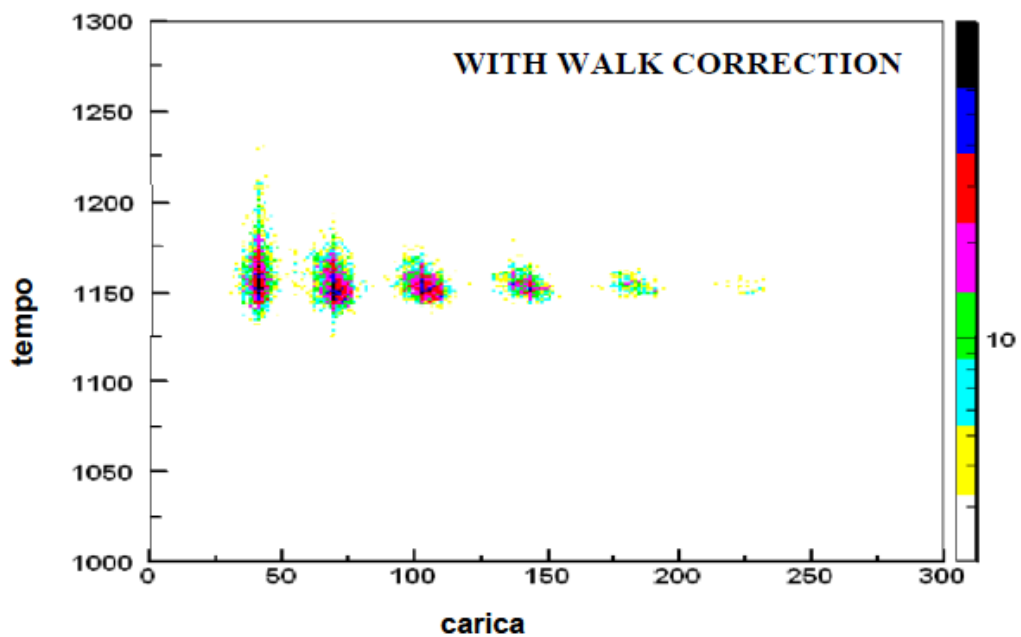


Figura 7.9: Scatter plot tempo-carica con correzione del walk.

Applicando questa procedura per ognuno dei tre SiPM presi in esame, nonché al valore di contributo dell'elettronica al rumore, si ottengono i seguenti valori.

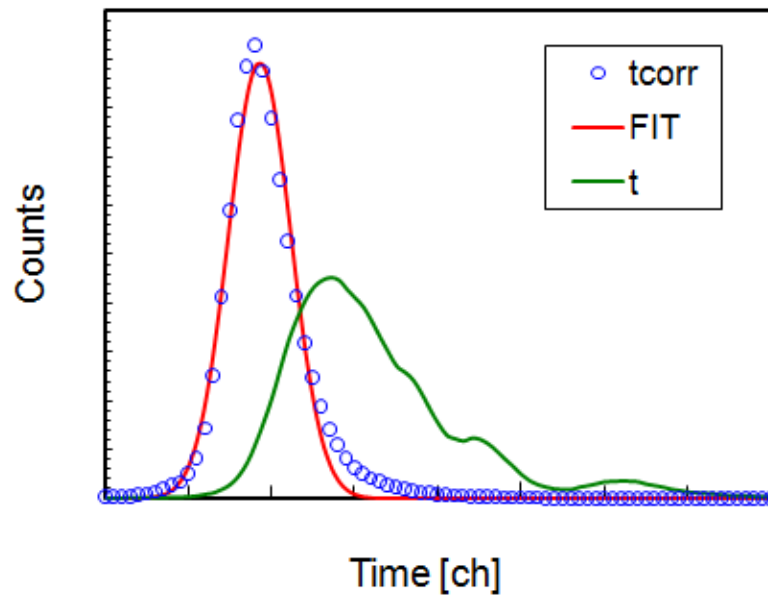


Figura 7.10: Confronto tra tempo prima e dopo la correzione del walk per uno dei sensori (Hamamatsu).

	Hamamatsu	SensL	STM	Rumore elettronico
FWHM	71,1 ns	189 ns	175 ns	68,4 ns

La scelta per le misure di eventi in coincidenza con i cristalli e la sorgente gamma cade ovviamente sul sensore Hamamatsu, che presenta la migliore risoluzione temporale.

Capitolo VIII:

Misure in coincidenza

Il lavoro svolto si colloca nell'ambito di TOPEM, un progetto che l'INFN sviluppa in collaborazione nazionale con diversi enti (tra cui l'ISS, Istituto Superiore di Sanità), volto alla realizzazione di una sonda per TOF-PET prostatica.

Il concept del progetto prevede l'utilizzo di un'array di SiPM accoppiati con cristalli che lavorano in coincidenza tra loro.

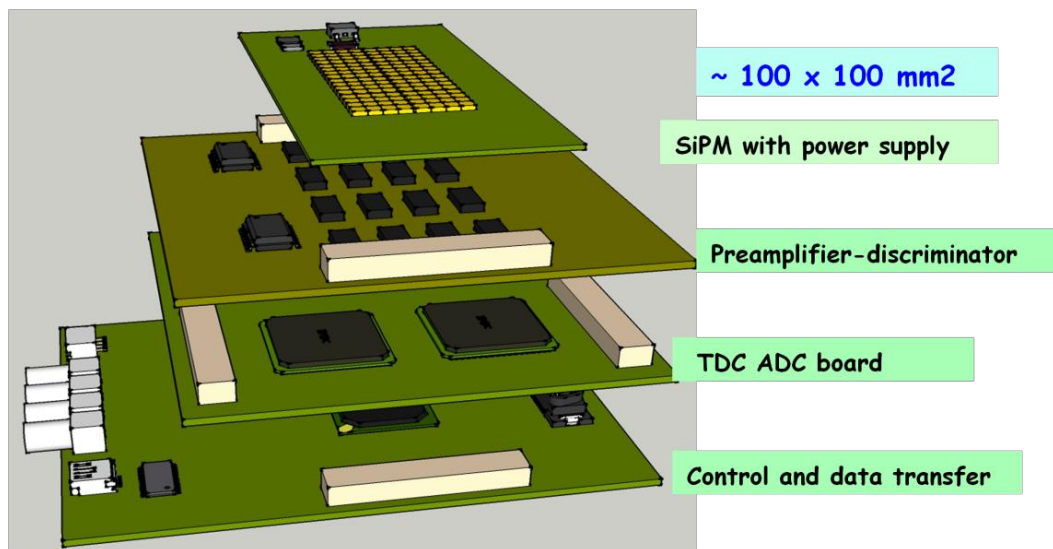


Figura 8.1: Rappresentazione schematica del prototipo di sonda.

Il primo passo verso la realizzazione del prototipo consiste quindi nell'andare ad esaminare il comportamento in coincidenza di due unità rivelatrici (SiPM e scintillatore) in presenza di una sorgente gamma, in particolare una sorgente ^{22}Na da 3.7kBq.

Il SiPM scelto presenta le seguenti dimensioni (in millimetri).

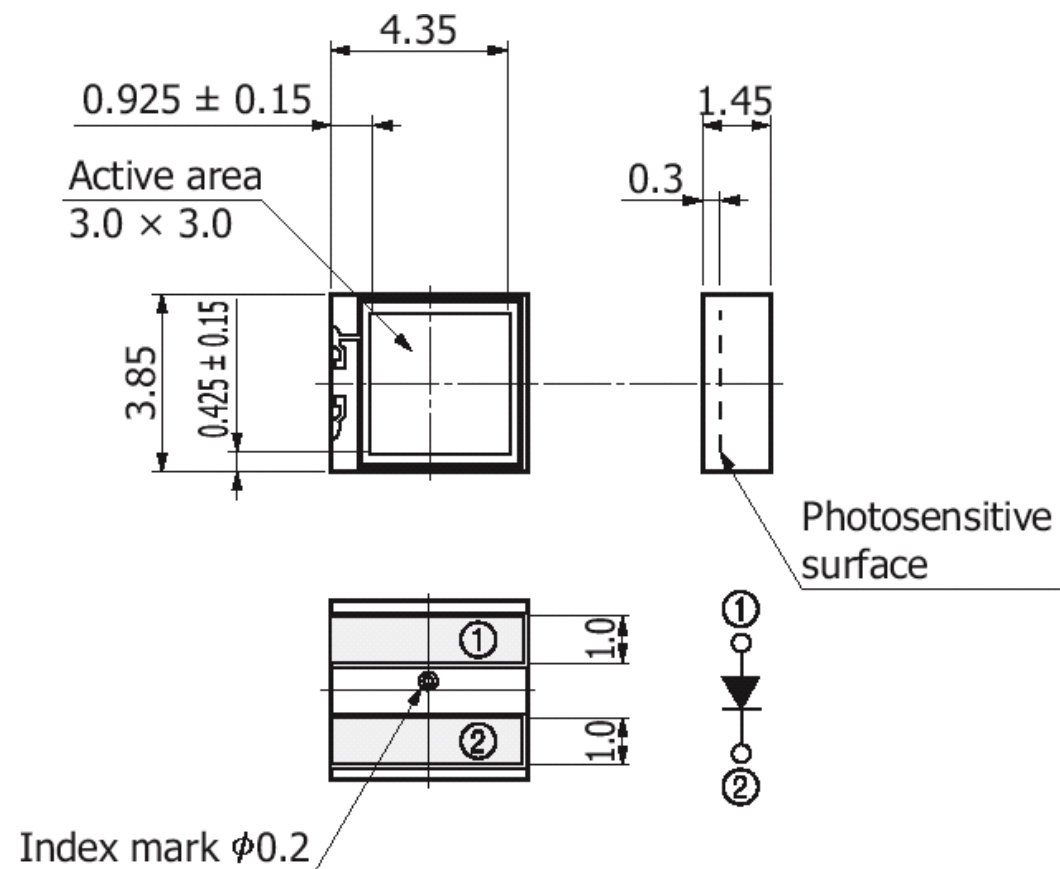


Figura 8.2: Sensore Hamamatsu in proiezioni quotate.

L'accoppiamento tra quest'ultimo e cristallo di LYSO rivestito è stato fatto in via provvisoria con dei supporti meccanici in PVC, di cui a seguire riportiamo le dimensioni.

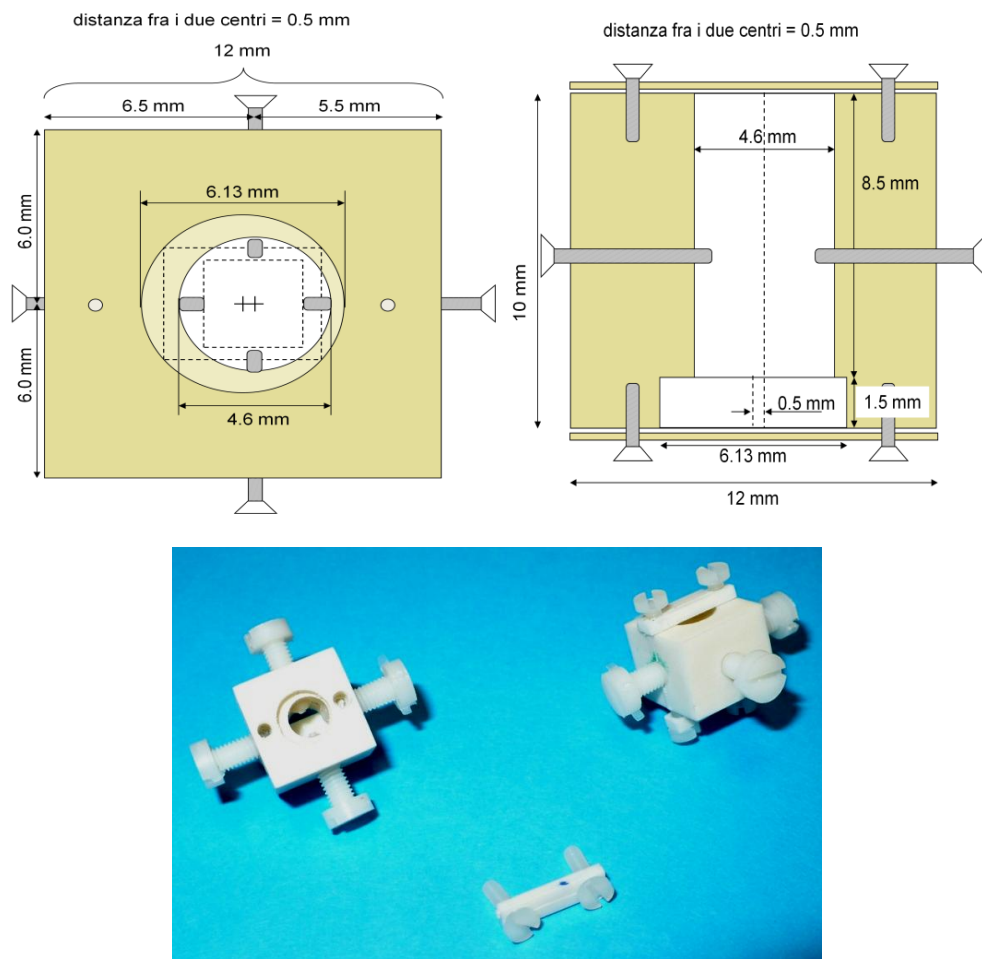


Figura 8.3: Proiezioni quotate del supporto meccanico (in alto), supporto meccanico finito (in basso).

In essi andrà inserito il cristallo di dimensioni $3 \times 3 \times 10 \text{ mm}^3$ che contatterà la faccia attiva del fotomoltiplicatore, saldato precedentemente su un supporto con il circuito di readout precedente, e infine le due unità rivelatrici andranno poste ad egual distanza dalla sorgente lungo il lato più lungo.

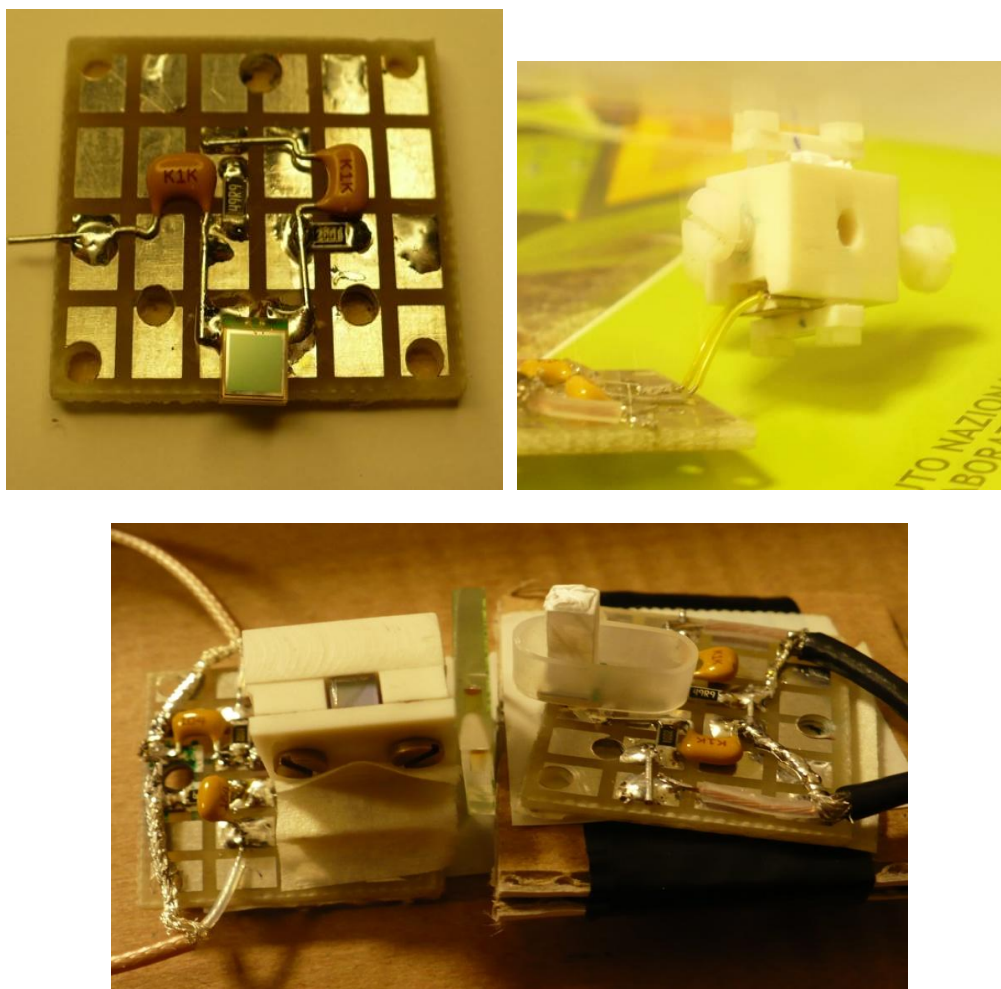


Figura 8.4: (dall'alto a sinistra in senso orario) SiPM saldato con circuito di polarizzazione, SiPM e supporto meccanico assemblati, due SiPM con sorgente radioattiva in mezzo.

Il setup sperimentale usato per acquisire gli spettri in coincidenza, in via molto semplificata, può essere riassunto nel seguente schema.

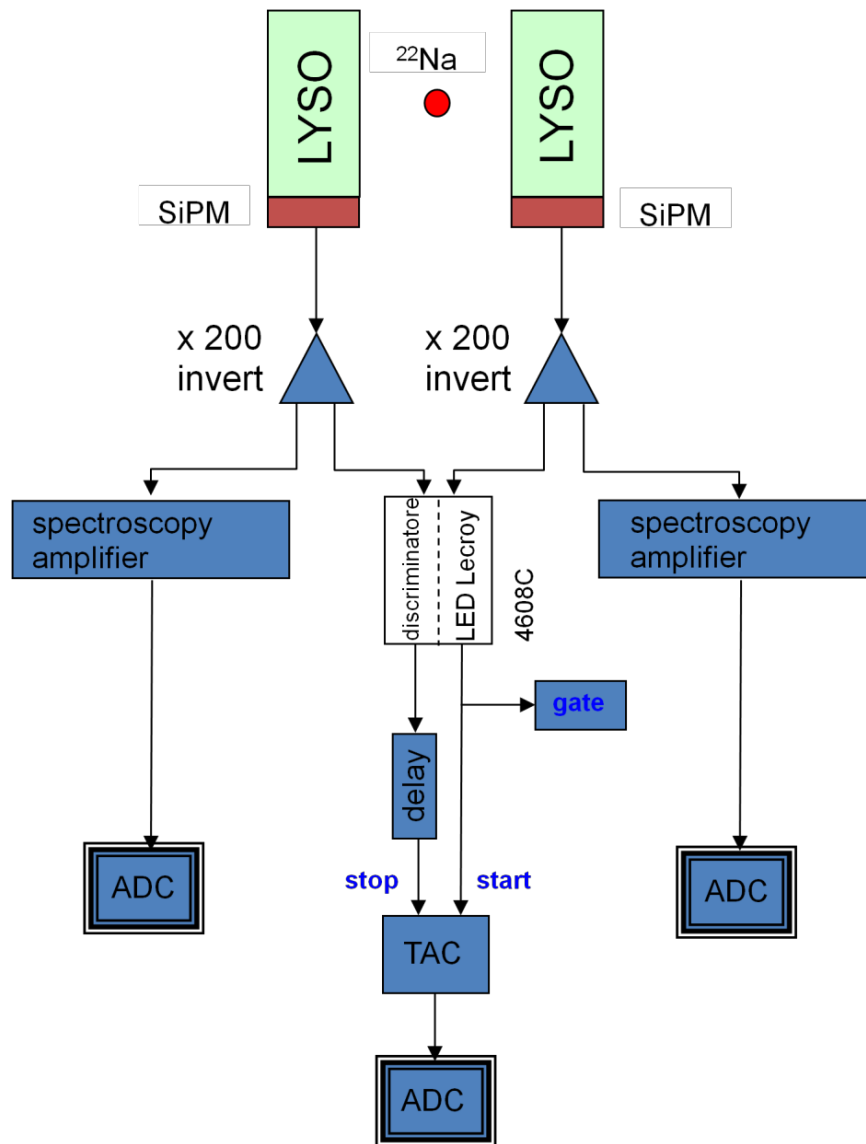
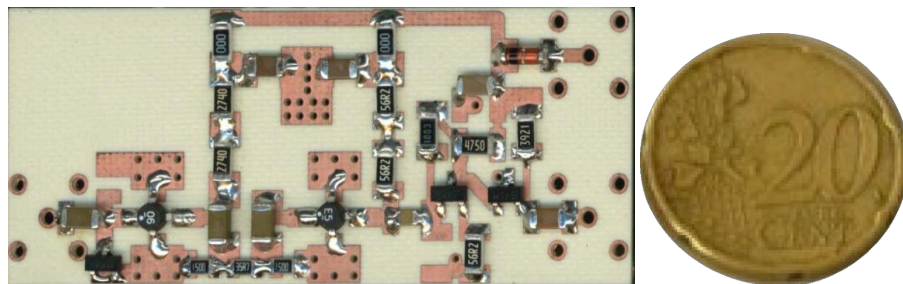


Figura 8.5: Schema elettronico per le misure in coincidenza.

Stavolta è cruciale la scelta di moduli che diminuiscano quanto più possibile il rumore introdotto, in quanto ci aspettiamo un peggioramento notevole della FWHM del picco di risoluzione temporale a causa della presenza di eventi di scintillazione non uniformi, al contrario della luce emessa dal laser, provenienti dai cristalli.

Sono state apportate quindi alcune modifiche rispetto all'apparato di acquisizione originario:

- L'amplificatore modulare Ortec è stato sostituito con dei pulse amplifier progettati ad hoc all'interno dei LNS.



Technical specifications:

Bandwidth: 1,5 KHz – 4Ghz;
Gain: 46 DB;
Noise figure: max 5 DB;
Input polarity: Positive;
Output polarity: Negative;
Power requirements: +12V 65mA.

Figura 8.6: Vista dall'alto dell'amplificatore (comparato in dimensioni con una moneta) e specifiche tecniche.

- Il linear FIFO è stato sostituito da un connettore a T, che non introduce rumore di fondo supplementare, ma dimezza l'ampiezza del segnale.
- Viene inserito anche un amplificatore spettroscopico programmabile a monte (C.A.E.N. N568B) dei canali dell'ADC predisposti all'acquisizione degli spettri energetici, al fine di avere un segnale che copre tutta la dinamica del convertitore avendo cura di evitare la

saturazione dei segnali (la tempistica di questi segnali è di qualche μs).

- Si ha lo stesso gate per tre segnali, ma per i segnali dei SiPM non è fondamentale stavolta che il gate sia ben centrato poiché l'ADC considera solo il valore massimo che si presenta all'ingresso nell'intervallo di acquisizione.

In prima battuta si procede dunque all'acquisizione degli spettri in energia dei due rivelatori e dello spettro tempo in coincidenza.

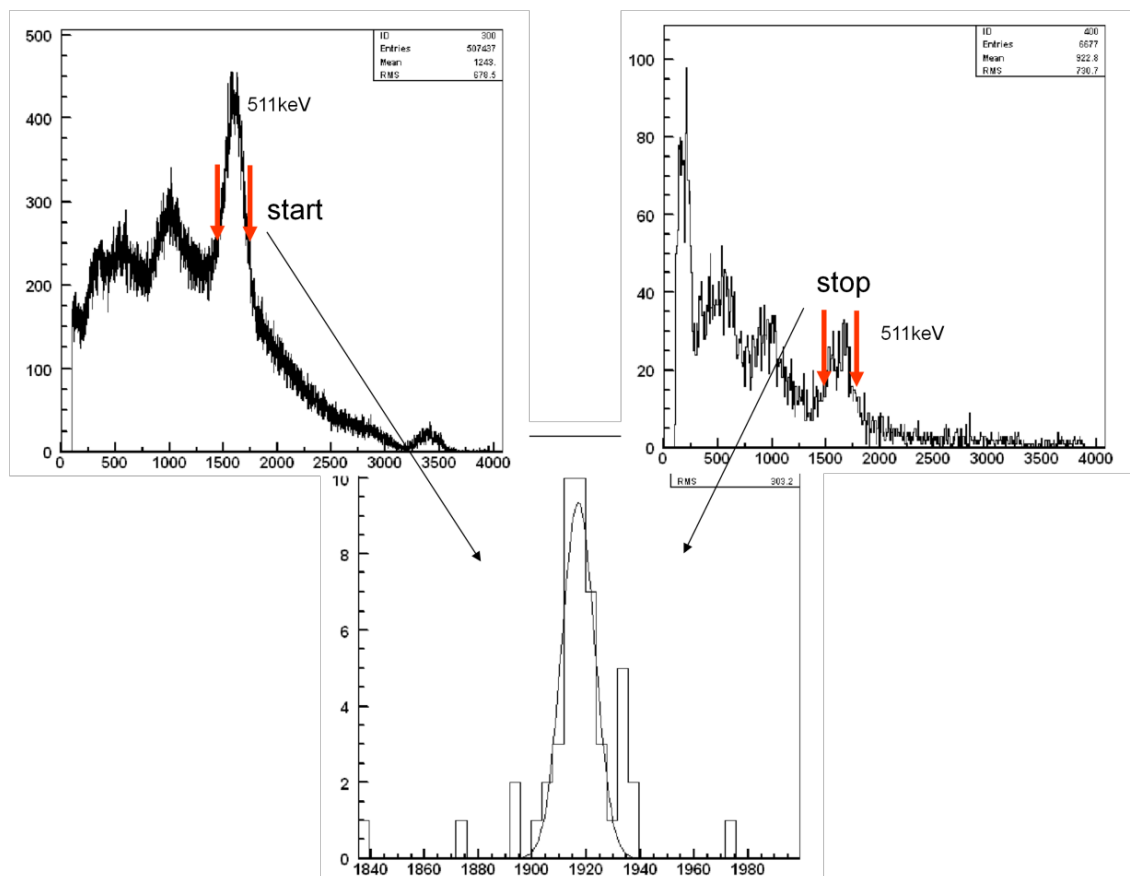


Figura 8.7: (dall'alto) spettro energetico di ciascuno dei due SiPM, spettro di coincidenza temporale.

Come previsto dalla teoria dell'interazione dei raggi gamma con la materia, lo spettro in energia presenta il picco fotoelettrico a 511 keV, un *Compton continuum* ad energie minori caratterizzato dalla sua spalla Compton (il picco precedente a 341 keV) e infine il picco di backscatter.

Si ottiene una risoluzione temporale FWHM di circa 350 ps, che è quasi la metà di quella dei sistemi TOF-PET attualmente in commercio.

Per avere una verifica sulla correttezza dei dati acquisiti basta vedere se il numero dei fotoni rilevati è congruente, almeno come ordine di grandezza, con il numero previsto da calcoli analitici:

$$\begin{aligned} \#ph &= (light\ yield) \cdot E \approx 30\ kph/MeV \cdot 0.511\ MeV \\ &\approx 15\ kph \end{aligned}$$

Per ricavare tale numero per via sperimentale innanzitutto bisogna ricavare il numero di celle eccitate dalla radiazione (*fired pixel*), il quale si ottiene banalmente come rapporto tra la carica associata al picco a 511 keV e quella del singolo fotone.

Attualmente siamo in regime di alta amplificazione quindi visualizziamo, per il SiPM che fornisce il trigger, solamente il picco da 511 keV, ma inserendo un attenuatore modulare tra i due amplificatori ci è possibile risalire al singolo fotone e fare il rapporto tenendo conto della differenza di amplificazione.

$$\#fired\ pixel = \frac{Q[511\ keV]}{Q[1\ fotone]} \approx 1400$$

Da questo numero è immediato ricavare il numero orientativo di fotoni, in quanto:

$$\#fotoni\ incidenti = \frac{\#fired\ pixel}{PDE} \approx 3000$$

Il numero ottenuto è ben cinque volte inferiore a quello atteso e, andando ad investigare il problema, se le misure sono state effettuate correttamente, l'unica motivazione possibile è che il valore di PDE (parametro che discuteremo nel prossimo capitolo) fornitoci dal costruttore del SiPM è falso, o meglio non tiene conto di molti effetti indesiderati che noi abbiamo provveduto ad eliminare dal computo, primo tra tutti il crosstalk.

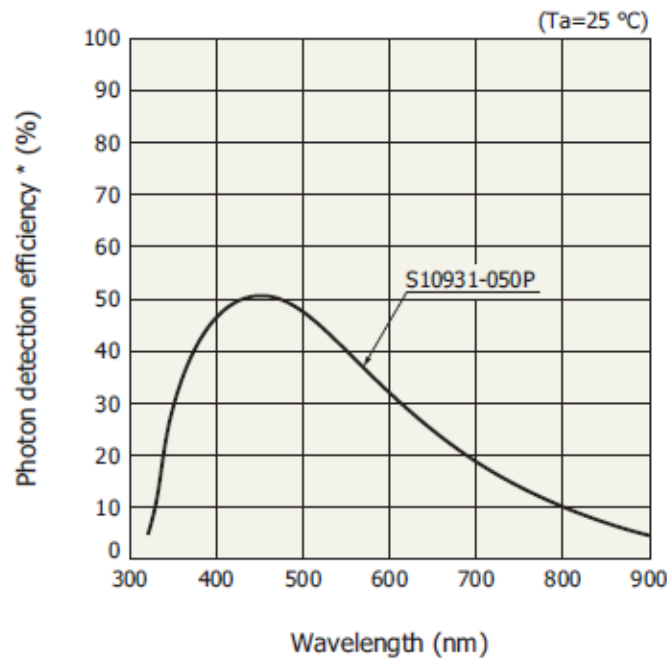


Figura 8.8: PDE in funzione della lunghezza d'onda secondo il datasheet del SiPM.

Capitolo IX:

Misure di PDE

Non tutti i fotoni che arrivano al fotosensore vengono rivelati, una grossa parte di questi pur incidendo su di esso non riesce a creare la valanga e quindi un segnale rivelabile.

Le principali ragioni che causano ciò si spiegano con l'**efficienza geometrica** ϵ_{geom} , determinata dalla dimensione relativa della regione inattiva fra le celle, l'**efficienza elettronica** ϵ_{el} , in quanto affinché sia possibile innescare la valanga occorre che la conversione del fotone avvenga in regioni dove il campo elettrico lo permetta, e inoltre per motivi fisici quali riflessione o assorbimento tra due strati successivi, rappresentati dall'**efficienza quantica QE**. L'efficienza complessiva del sensore viene generalmente definita come **PDE** (*photon detection efficiency*) e ad essa è legato il numero effettivo di fotoni rivelati dal sensore.

$$PDE = \epsilon_{geom} \cdot \epsilon_{el} \cdot QE$$

L'apparato strumentale per la misura della PDE è il seguente:

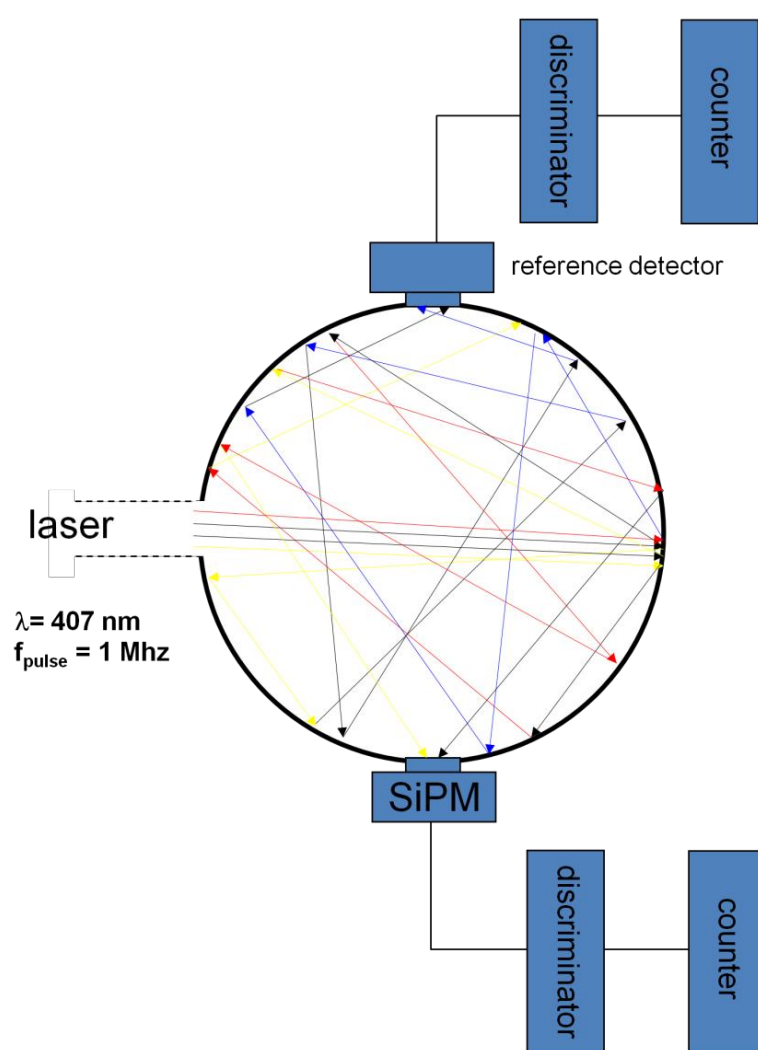


Figura 9.1: Schema elettronico per le misure di PDE.

In ordine da valle a monte abbiamo:

- *Contatore FPGA*

Si tratta di un sistema basato su FPGA, sviluppato appositamente per il conteggio di impulsi al secondo inviati dai SiPM, realizzato per un progetto DMNR, parallelo a TOPEM, e di cui parleremo successivamente.

- *Discriminatore*

Il modulo è lo stesso di quello utilizzato nelle precedenti misure, con la differenza che stavolta viene usato impostando la soglia sia a 0.5 fotoni che a 1.5 fotoni.

Questo perché vogliamo essere sicuri di lavorare in un regime di bassa intensità laser cosicché si contino solo eventi di singolo fotone, quindi se il rapporto tra i conteggi a 0.5 e 1.5 fotoni rimane costante siamo sicuri di essere in tali condizioni.

- *Sfera di Ulbricht*

Detta più comunemente sfera integratrice, è un dispositivo per l'appunto sferico e cavo, il cui interno è rivestito di una vernice speciale bianca e diffondente (Lambertiana) nello spettro del visibile.



Figura 9.2: Sfera integratrice.

La legge di Lambert ci dimostra che, a causa delle continue riflessioni, l'illuminamento di ogni punto della superficie interna della sfera è costante e proporzionale al flusso totale emesso dal laser.

Sia r la distanza tra una sorgente puntiforme S e una porzione di superficie $\Delta A'$ orientata; la proiezione di $\Delta A'$ sopra la superficie sferica di centro S e raggio r è:

$$\Delta A = \Delta A' \cdot \cos \alpha$$

Dove α è l'angolo compreso tra le due normali a $\Delta A'$ e ΔA .

L'angolo solido sotto cui $\Delta A'$ è vista da S risulta quindi:

$$\Delta \Omega = \frac{\Delta A}{r^2} = \frac{\Delta A' \cdot \cos \alpha}{r^2}$$

Il flusso di radiazione emesso entro l'angolo solido $\Delta \Omega$ è:

$$\Delta \Phi = I \cdot \Delta \Omega = I \cdot \frac{\Delta A' \cdot \cos \alpha}{r^2}$$

Concludendo, l'irraggiamento $E = \Delta \Phi / \Delta A'$ sopra la superficie sferica A' è:

$$E = \frac{I \cos \alpha}{r^2}$$

Nel caso in cui la radiazione colpisce perpendicolarmente la superficie, si avrà $\alpha = 0$, quindi la formula diventa:

$$E = \frac{I}{r^2}$$

Quindi dipende solo dall'intensità del fascio laser e dalle dimensioni della sfera.

Possiamo dunque affermare che per i due fotosensori collocati sulle finestre della sfera, uno che funge da sensore di riferimento e di cui è nota la PDE per misure fatte all'INAF di Catania e già pubblicate e l'altro di cui vogliamo misurarla, si ha:

$$(\#fotoni)_{ref} = (\#fotoni)_x$$

Dato che vale inoltre:

$$\#counts = PDE \cdot Area \cdot \#fotoni$$

Otteniamo l'espressione finale della PDE:

$$PDE_x = PDE_{ref} \cdot \frac{A_{ref}}{A_x} \cdot \frac{(\#counts)_x}{(\#counts)_{ref}}$$

Come sensore di riferimento è stato usato quello STM, quindi il rapporto delle aree è 1/36, mentre la PDE, come detto, proviene da dati sperimentali acquisiti per esperimenti precedenti ed è 7.3%.

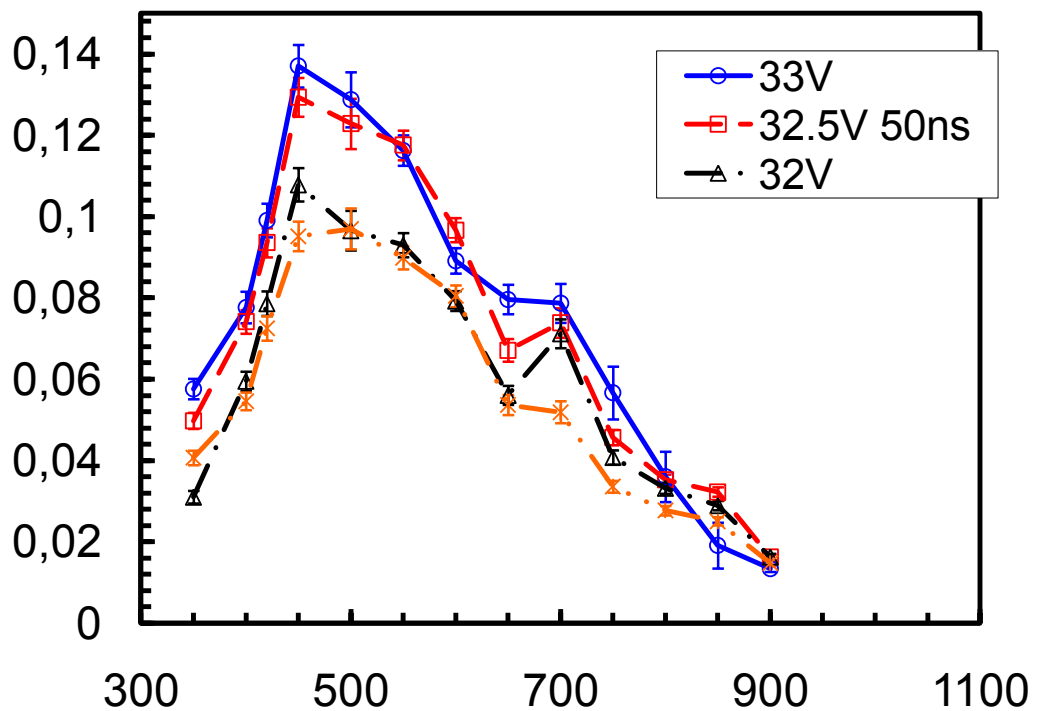


Figura 9.3: Misure di PDE per il sensore STM effettuate all'INAF.

Si procede dunque all'acquisizione dei dati relativi ai conteggi al secondo, ma, data la variabilità del dark count con la temperatura, si è escogitato un modo per avere una misura assoluta.

In sostanza non si fa altro che acquisire per in successione 10 minuti con il laser spento, 10 minuti con l'esposizione del laser e infine altri 10 minuti con laser spento.

Così facendo possiamo interpolare ai minimi quadrati una retta che ci permetta di sottrarre i conteggi di buio a quelli complessivi.

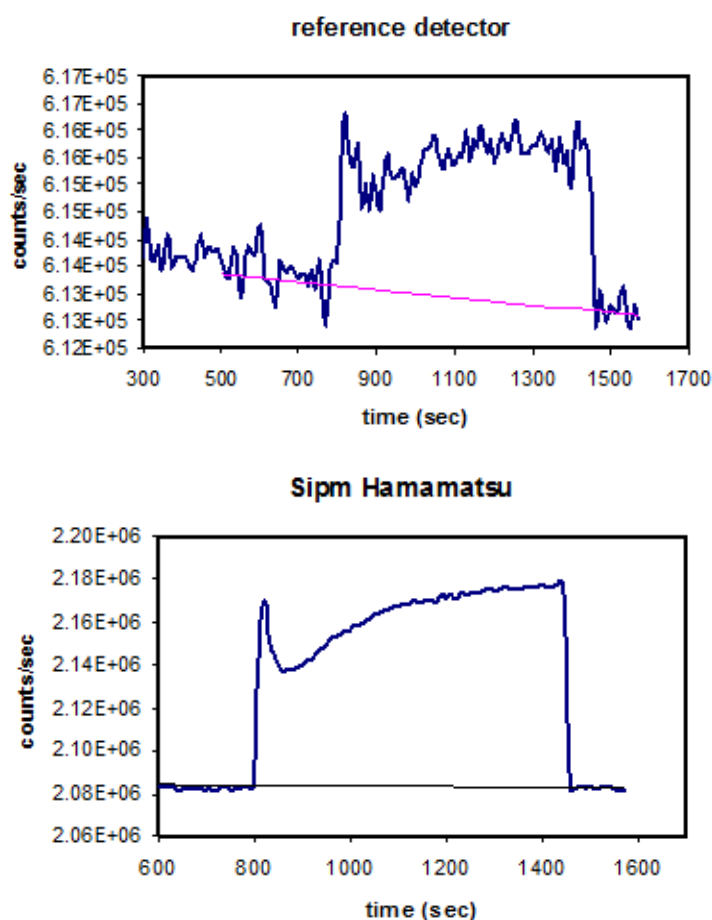


Figura 9.4: variazione dei conteggi al secondo nel tempo e in presenza del laser per SiPM di riferimento (in alto) e Hamamatsu (in basso).

Gli andamenti temporali evidenziano una non linearità nell'andamento del laser, ma poco importa in quanto tale fenomeno si riflette su entrambi i SiPM e quindi non inficia sulla misura finale.

Si sottolinea inoltre la difficoltà nel trovare un'intensità luminosa ottimale per due SiPM diversi, in quanto per quello STM si è a limite con il non rivelare nulla, mentre per il SiPM Hamamatsu c'è il rischio di avere eccessiva esposizione e quindi *double firing*; da

qui l'importanza di controllare sempre i rapporti tra i conteggi a 0.5 e 1.5 fotoni.

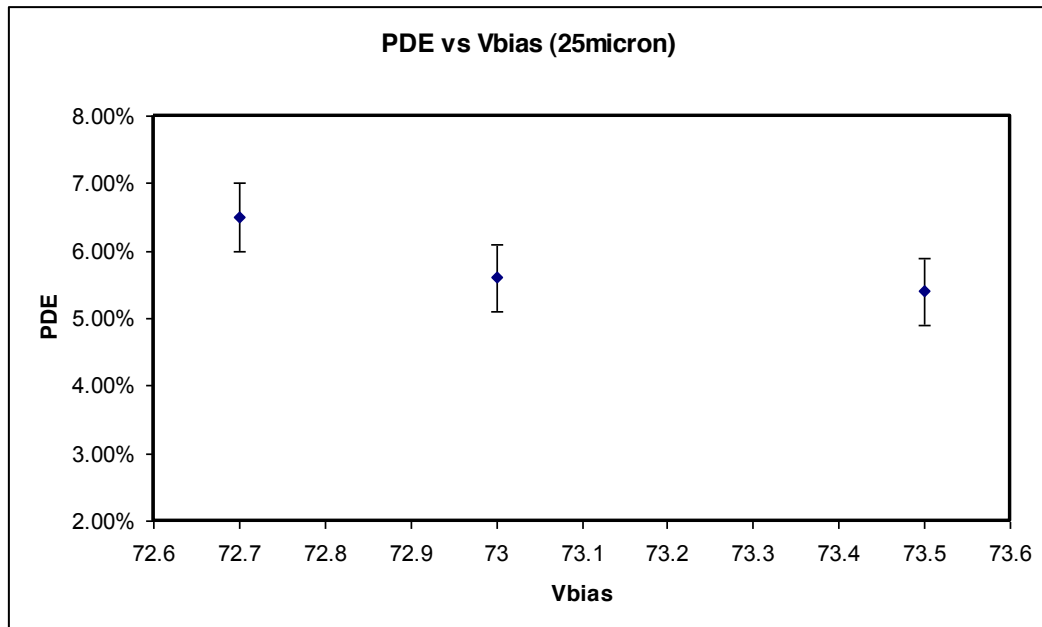


Figura 9.5: PDE al variare dell'overvoltage per Hamamatsu da 25 μm .

Inserendo il valore di PDE così ottenuto nella formula del capitolo precedente per il conteggio dei fotoni incidenti, si ottiene un valore di circa 20000 fotoni, che non coincide esattamente con i 15000 del calcolo teorico, ma, considerando che siamo sullo stesso ordine di grandezza e che quella che facciamo più che un calcolo è una stima del numero di fotoni, è un risultato più che soddisfacente.

In seconda battuta oltre al sensore Hamamatsu da 25 μm è stata misurata anche la PDE di un altro SiPM, sempre Hamamatsu, con cella da 50 μm .

Quello che ci aspettiamo è, dato che il dispositivo da 50 μm ha un filla factor doppio dell'altro (61.5% contro 30.8%), una PDE doppia per questo nuovo sensore.

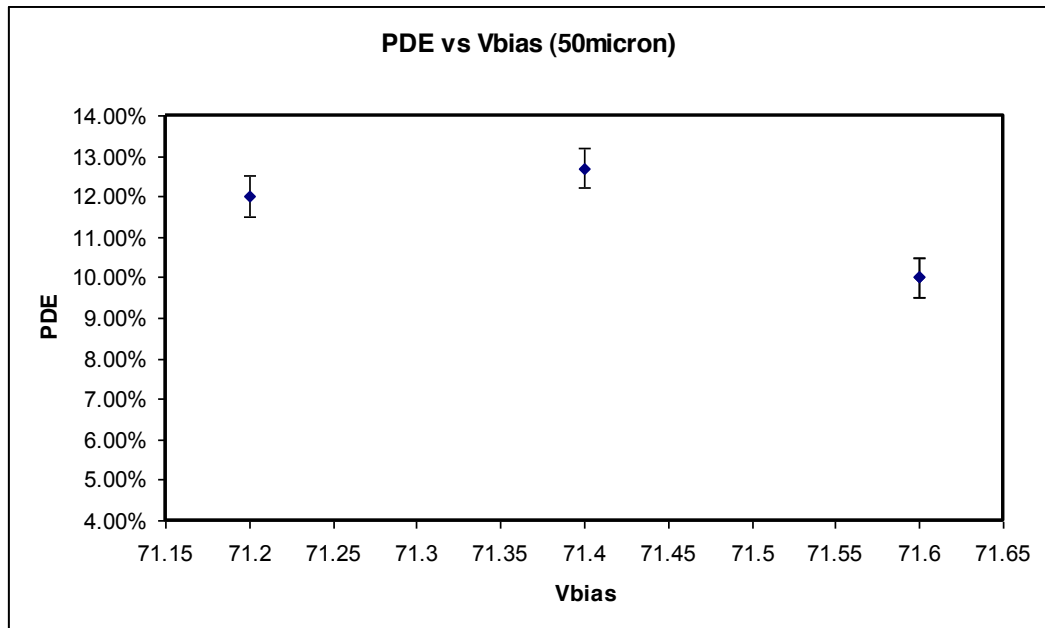


Figura 9.6: PDE al variare dell'overvoltage per Hamamatsu da 50 μm .

Come volevasi dimostrare si ottiene proprio quanto atteso, ciò è un risultato molto importante anche dal punto di vista della risoluzione temporale, in quanto si ha che quest'ultima segue una legge di proporzionalità:

$$\propto \frac{1}{\sqrt{PDE_2/PDE_1}}$$

Dunque con una PDE maggiore si ha un miglioramento delle prestazioni temporali e quindi una maggior precisione spaziale nella

rivelazione, in particolare se sostituissimo il sensore Hamamatsu nelle misura di timing otterremmo una risoluzione di:

$$\frac{350}{\sqrt{2}} \approx 250 \text{ ps}$$

Valore che è in linea con quanto si trova in letteratura scientifica sulla misura di timing con Hamamatsu da 50 μm .

Capitolo X:

Analisi di dark count

Lo scopo della caratterizzazione è studiare il comportamento della vasta gamma di SiPM a disposizione dal punto di vista di *crosstalk* e *dark noise*.

Tali fenomeni sono entrambi sorgenti di rumore, l'uno ottico ed elettrico e l'altro termico, e sono già stati introdotti nel capitolo riguardante i SiPM, ma verranno ampiamente approfonditi a seguire. Abbiamo precedentemente definito il dark count come il numero di impulsi per unità di tempo in assenza di luce; nel silicio infatti diversi processi statistici di natura termica causano la produzione spontanea di portatori di carica.

Se questi portatori, generati in modo casuale, arrivano alla zona di moltiplicazione, possono dare origine a una scarica non distinguibile da quelle prodotte da eventi reali.

Ciò significa che il segnale risultante da un fotoelettrone generato nella regione di svuotamento è lo stesso di quello prodotto da un portatore generato casualmente.

Il processo fisico che causa il dark rate è intrinseco alla struttura a bande del silicio: dovendo sempre verificarsi la cosiddetta **legge di azione di massa**

$$n p = n_i p_i$$

dato che un semiconduttore drogato deve sempre mantenere una carica neutra, in esso avviene una continua produzione e ricombinazione di portatori di carica ed è proprio questa produzione la responsabile della corrente di buio.

Una formula per il rate di generazione nel silicio che tiene in considerazione i vari processi che la causano è quella data dall'equazione di Shockley-Read-Hall

$$G = \frac{n_i^2 - p n}{\tau_{e0} \left(p + n_i e^{-\frac{E_T - E_0}{KT}} \right) + \tau_{h0} \left(p + n_i e^{-\frac{E_T - E_0}{KT}} \right)}$$

dove E_T ed E_0 sono rispettivamente i livelli energetici dell'impurezza e del livello di Fermi, mentre τ_{e0} e τ_{h0} sono i tempi caratteristici di cattura e rilascio nei centri di intrappolamento di elettroni e lacune.

G rappresenta il rate netto di generazione, che diventa negativo se la ricombinazione supera la generazione.

Questi portatori così generati sono quelli che, se diffusi nella zona di svuotamento, possono causare la valanga che rappresenta il dark count.

Dall'equazione si evince inoltre un ruolo fondamentale della temperatura: maggiore è la temperatura, maggiore sarà la produzione di cariche per effetto termico.

Il cross talk invece è caratteristico di tutti i dispositivi a matrice e si presenta quando due o più pixel interagiscono tra loro o per ragioni ottiche (cross talk ottico) o per cause di natura elettronica (cross talk elettronico).

Il cross talk ottico si verifica quando, durante una fotorivelazione, i portatori che formano la corrente inversa dovuta alla valanga generata dal fotone incidente, emettono dei fotoni che vengono detti di Bremsstrahlung e che si propagano lungo il silicio come su delle guide d'onda.

Se uno di questi fotoni, dovuti proprio al meccanismo interno di funzionamento dello stesso dispositivo, raggiunge l'area attiva di un altro pixel del sensore e viene assorbito, innescando in questo una moltiplicazione a valanga, viene generato un impulso spurio fortemente correlato all'assorbimento del fotone primario; tale impulso si presenta come il contributo del cross talk al rumore del sensore.

La relazione che si ha fra la corrente inversa e la produzione interna di fotoni, intesa come flusso, è data dalla seguente equazione

$$I = \frac{I_{reverse}}{q} \cdot \gamma$$

dove $I_{reverse}$ è la corrente inversa che attraversa la giunzione, q è la carica dell'elettrone e γ il coefficiente di emissione.

Nonostante γ sia molto piccolo, cioè si ha una bassa probabilità di emissione di fotoni (circa un fotone ogni 10^5 elettroni), la probabilità di avere cross talk ottico è significativa, ciò si deve alla notevole sensibilità degli SPAD che compongono il SiPM.

Il segnale di ogni singolo pixel può venire contaminato dai fotoni derivanti dalle microcelle vicine anche in funzione del suo coefficiente di assorbimento α , che, come per tutti i fotorivelatori a stato solido, è strettamente legato alla lunghezza d'onda λ della luce in questione.

$$I_{\text{photon}}(\lambda, d) = I_{\text{photon}}(\lambda, 0)e^{-\alpha(\lambda)d}$$

Inoltre l'attenuazione di tale flusso di fotoni dipende ovviamente anche dal percorso nel dispositivo.

Il cross talk elettronico invece ha luogo quando dei portatori di carica, emessi dalla giunzione di un pixel che ha assorbito un fotone, diffondono attraverso la regione epitassiale di tipo p^+ comune a tutte le microcelle del SiPM.

Qui possono essere assorbiti da pixel vicini, dove innescano valanghe che introducono impulsi spuri scorrelati dalla reale rivelazione di fotoni (anche per questo processo l'elevata sensibilità del sensore non gioca un ruolo del tutto favorevole).

I sistemi utilizzati nei SiPM per attenuare il fenomeno del cross talk sono principalmente due: aumentare la minima distanza fra le zone attive di due pixel adiacenti, il cosiddetto *pitch*, o a realizzare uno

scavo fra un pixel e un altro, delle *trench* riempite di ossido, così da realizzare un isolamento ottico-elettrico tra le varie microcelle.

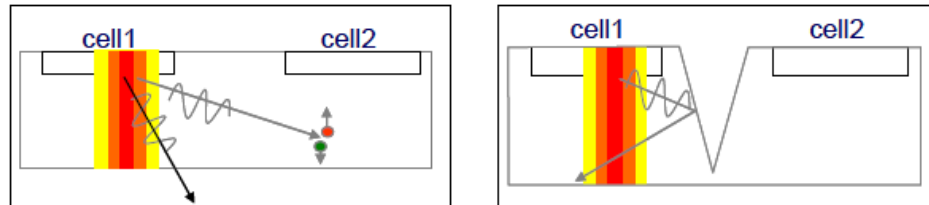


Figura 10.1: Benefici del trench.

Tra le due soluzioni la prima ridurrebbe di gran lunga l'area sensibile del dispositivo, diminuendo così l'efficienza geometrica e quindi l'intera efficienza di raccolta del SiPM; quindi si preferisce il secondo metodo con la realizzazione di trench molto sottili che non intaccano eccessivamente il fill factor del sensore.

Il setup sperimentale utilizzato per misurare il dark count e risalire alla probabilità di cross talk è il seguente.

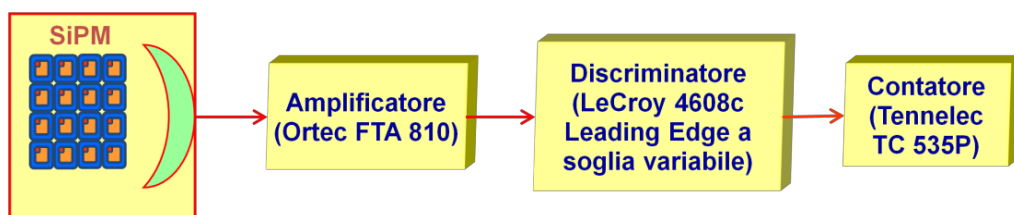


Figura 10.2: Schema elettronico per il dark count.

Le misure effettuate prevedono la rilevazione dei valori di dark al variare della soglia del discriminatore, in modo da tracciare in

prima battuta delle curve di dark count e normalizzarle al valore di rate massimo e al valore di soglia corrispondente al singolo fotone (è il valore al quale il rate si dimezza).

Questo lavoro è stato fatto per numerose tipologie di sensori di diversi produttori al fine di avere una valutazione comparata in base al dark count.

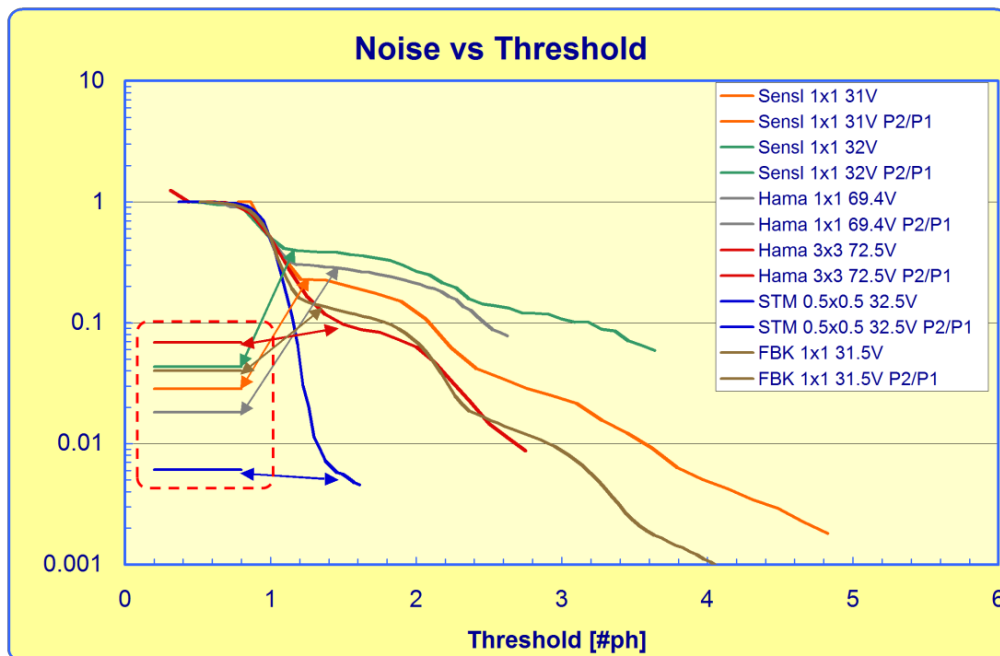


Figura 10.3: Curve di dark count per diversi SiPM.

Gli andamenti delle curve evidenziano delle zone a pendenza minore, detti *plateau*, in cui il dark count è idealmente costante e i quali rappresentano i valori di soglia corrispondenti a $n/2$ fotoni, invece le zone a pendenza più ripida cadono intorno ai valori di soglia corrispondenti a n fotoni ($n \in \mathbb{N}$).

In figura, all'interno del riquadro tratteggiato, è indicato il valore teorico di rumore in assenza di crosstalk, stimato secondo statistica di Poisson.

Sempre ricorrendo ad alcuni concetti di statistica di Poisson è possibile ricavare dalle curve di dark count i valori di crosstalk.

Ci aspettiamo eventi singoli con probabilità per unità di tempo μ piccola (in realtà non sarà per unità di tempo ma sarà normalizzata alla durata del segnale logico del contatore, che è di qualche decina di nsec tipicamente).

$$P(\mu, k) = \frac{\mu^k}{k!} e^{-\mu}$$

Facendo il rapporto tra la probabilità di avere due eventi e la probabilità di avere un evento otteniamo

$$\frac{P(\mu, 2)}{P(\mu, 1)} = \frac{\mu}{2}$$

In prima approssimazione (rough) questo rapporto si può calcolare come rapporto tra il tasso di conteggi con soglia a due celle e quello con soglia a una cella; in realtà con soglia ad una cella si includono anche tutti i successivi valori di probabilità, così come con soglia a due celle.

Con soglia al minimo si rivelano tutti gli eventi con almeno una cella attiva.

Poichè nella distribuzione di Poisson viene contemplato anche il caso di zero celle, allora il numero di conteggi con soglia al minimo, normalizzato a 1, sarà

$$N_{MinTh} = 1 - P(\mu, 0) = 1 - e^{-\mu}$$

da cui ci può ricavare μ

$$\mu = -\ln(1 - N_{MinTh})$$

Ciò è valido sotto le ipotesi di distribuzione Poissoniana, cioè eventi indipendenti con probabilità molto piccola, questo si verifica se guardiamo il fenomeno su una scala temporale opportuna (dell'ordine dei nsec).

N_{MinTh} si ricava banalmente moltiplicando la durata del segnale logico del contatore per il rate totale al secondo.

Combinando questa equazione con la precedente otteniamo

$$\frac{P(\mu, 2)}{P(\mu, 1)} = \frac{\mu_{uniform}}{2} = -\frac{1}{2} \ln(1 - N_{MinTh})$$

Ossia il valore di $\mu/2$ dedotto nell'ipotesi micro-poissoniana: gli eventi sono distribuiti uniformemente nel tempo e dunque dividendo il tasso misurato per la durata si ottiene il tasso microscopico (per nanosecondo).

Se quest'ipotesi non è vera (ad esempio c'è cross talk), allora gli eventi non saranno distribuiti uniformemente nel tempo, ma ci sarà correlazione.

Per verificare ciò basta esaminare la distribuzione misurata a vari valori di soglia: denominati S_1 il tasso misurato con soglia a 0.5 fotoni, S_2 con soglia a 1.5 fotoni, S_3 con soglia a 2.5 fotoni si ha che

$$\frac{S_2 - S_3}{S_1 - S_2} = \frac{P(\mu, 2)}{P(\mu, 1)} = \frac{\mu_{macro}}{2}$$

ma in questo caso, poiché stiamo lavorando con la probabilità integrata nel tempo, la non è necessariamente relativa alla micro-distribuzione, o meglio, se otteniamo lo stesso valore significa che la macro-distribuzione e la micro-distribuzione sono identiche.

Se invece si ottiene un valore diverso, significa che ci sono degli eventi correlati, ed è proprio questo quello che accade realmente e che ci permette di misurare il cross talk.

Adesso, nel caso della macro-correlazione, si ha che

$$\begin{aligned} P(\mu, k > 1) &= 1 - P(\mu, 0) - P(\mu, 1) = 1 - e^{-\mu} - \mu e^{-\mu} \\ &= 1 - (1 + \mu)e^{-\mu} \end{aligned}$$

che rappresenta la probabilità di avere più celle attive simultaneamente.

La stessa formula, utilizzando la relativa alla micro-distribuzione moltiplicata per la durata del segnale logico, ci fornisce la probabilità di avere più celle attive nell'ipotesi di uniformità temporale (dunque senza correlazione).

La differenza tra le due probabilità ci darà la probabilità di avere più celle attive a causa di cross talk

$$P(\mu_{macro}, k > 1) - P(\mu_{micro}, k > 1) = P_{crosstalk}$$

Si ricavano così i valori di cross talk per tutti i modelli di SiPM.

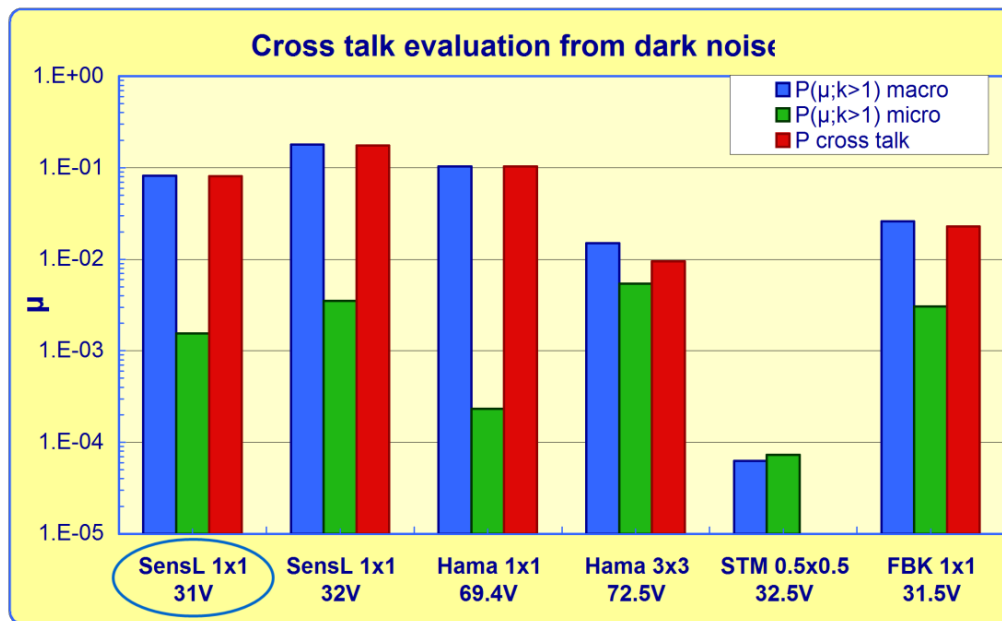


Figura 10.4: Valori di cross talk per i vari SiPM.

Inoltre, avendo a disposizione ben 24 SiPM della SensL da 1mmx1mm, sono stati caratterizzati per individuare eventuali differenze significative tra sensori nominalmente identici.

Il risultato di tale studio è una buona omogeneità tra i sensori.

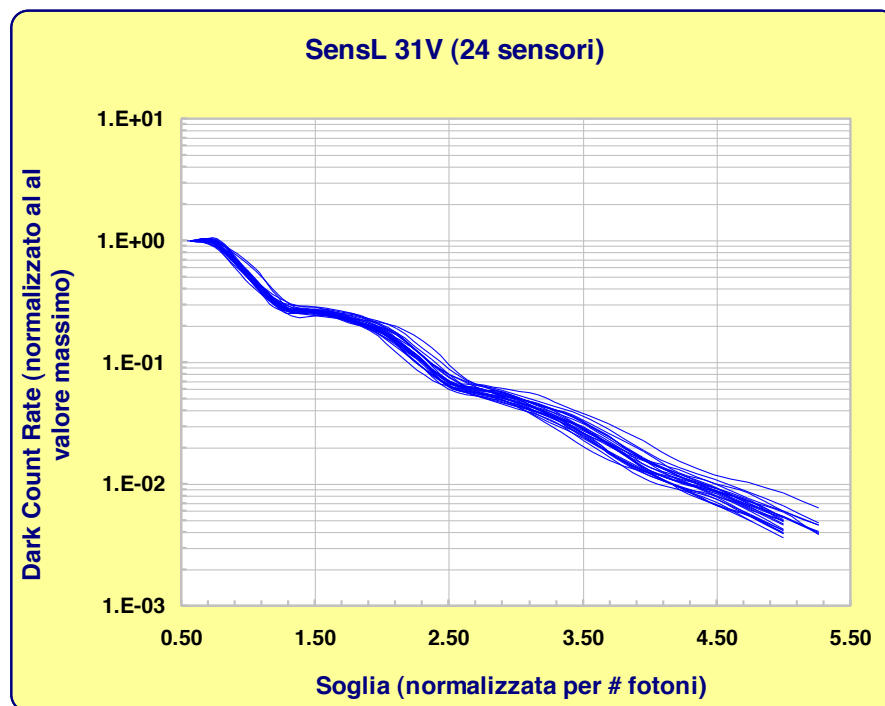


Figura 10.5: Spread delle curve di dark count per i SiPM SensL.

Capitolo XI:

Conclusioni

Tirando le somme, sono stati testati con successo i rivelatori di radiazione nucleare per TOF-PET costituiti da SiPM e cristallo scintillatore.

I risultati prodotti dal punto di vista della risoluzione temporale, per il progetto TOPEM, sono di circa 350 psec per gli Hamamatsu con cella da 25 μm e di circa 250 psec (teorizzati) per gli Hamamatsu da 50 μm .

Traducendo la risoluzione temporale in risoluzione spaziale otteniamo rispettivamente 5.25 cm e 3.75 cm.

Le prospettive di ricerca nel settore della TOF-PET, però, si propongono di diminuire ulteriormente di questi valori mediante l'utilizzo di nuovi materiali scintillanti e SiPM che permettano un timing più veloce.

Bibliografia

- [1] “Caratterizzazione di dispositivi fotomoltiplicatori al silicio per applicazioni di monitoraggio di radiazioni”, G. Greco
Master Thesis, Università degli Studi di Catania.
- [2] “Studio di fotomoltiplicatori al silicio (SiPM)”, M. Perani
Master Thesis, Università di Bologna.
- [3] “Radiation detection and measurement”, G.F. Knoll.
- [4] <http://en.wikipedia.org/>
- [5] “Ultra precise timing with SiPM-Based TOF PET scintillation detectors”, S. Seifert, R. Vinke, H. T. van Dam, H. Lohner, P. Dendoveen, F. J. Beekman, D. R. Schaart.
- [6] “Multi-pixel photon counters for TOF PET detector and its challenges”, C. L. Kim, G. C. Wang, S. Dolinsky.
- [7] “Factor influencing time resolution of scintillators and ways to improve them”, P. Lecoq, E. Auffray, S. Brunner, H. Hillemanns, P. Jarron, A. Knapitsch, T. Meyer, F. Powolny.
- [8] “A measurement of the photon detection efficiency of silicon photomultipliers”, A. N. Otte, J. Hose, R. Mirzoyan, A. Romaszkievicz, M. Teshima, A. Thea.
- [9] “Presentazione TOPEM meeting Catania 2010”, F. Garibaldi.