



UNIVERSITÀ DEL SALENTO

FACOLTÀ DI SCIENZE MATEMATICHE FISICHE E NATURALI

Corso di Laurea Magistrale in Fisica

**TEST DI CHIP DI LETTURA
PER IL RIVELATORE A PIXEL DI ATLAS
AD HIGH-LUMINOSITY LHC**

CERN-THESIS-2020-372

15/10/2020

candidata:

Alessandra Palazzo



Relatori:

Prof.ssa Stefania Spagnolo

Prof. Enrico Junior Schioppa

Indice

Introduzione	i
1 La produzione di coppie di Higgs	1
1.1 Il bosone di Higgs	1
1.1.1 Canali di produzione e di decadimento	4
1.1.2 L'accoppiamento trilineare λ	6
1.2 Produzione di coppie di Higgs in collisioni pp	7
1.2.1 Gluon-gluon fusion	9
1.2.2 Vector boson fusion	9
1.3 La ricerca di coppie di Higgs ad LHC	11
1.3.1 Canale gluon-gluon fusion	11
1.3.2 Canale vector boson fusion	14
1.4 Interpretazione dei risultati: produzione non risonante	16
1.5 Interpretazione dei risultati: produzione risonante	18
2 Il b-tagging ad ATLAS	20
2.1 Segnatura dei b -jet	20
2.2 Oggetti fisici utilizzati per il b -tagging	22
2.3 Valutazione dell'efficienza degli algoritmi di b -tagging tramite simulazione Monte Carlo	23
2.4 Calibrazione degli algoritmi di b -tagging sui dati	23
2.5 Algoritmi per l'identificazione di b -jet	25
2.5.1 Algoritmi di primo livello	25
2.5.2 Algoritmo di secondo livello	29
2.6 Modifiche all'Inner Detector durante il run 2	30
2.6.1 Impatto sulla ricostruzione delle tracce	33
2.6.2 Impatto sull'efficienza di b -tagging	34
2.7 Inner Tracker per HL—LHC	36
2.7.1 Impatto sull'efficienza di b -tagging	38
2.8 Impatto del b -tagging nella ricerca di coppie di Higgs ad HL—LHC	42
2.9 Evoluzione futura degli algoritmi di b -tagging	43
3 Rivelatori a semiconduttore	47
3.1 Proprietà generali dei solidi	47
3.1.1 Conduttori, semiconduttori ed isolanti	48
3.2 Rivelatori al Silicio	49

3.2.1	Giunzione pn	50
3.3	Formazione del segnale	54
3.4	Danni da radiazione	55
3.5	Rivelatori a pixel	55
3.5.1	Rivelatori ibridi	56
3.5.2	Rivelatori planari e 3D	58
3.6	Confronto tra i rivelatori a pixel dell'Inner Detector e di ITk	60
3.6.1	Sensori 3D	61
3.6.2	Sensori planari	61
4	Il chip RD53A	63
4.1	Caratteristiche generali	63
4.2	Struttura del chip	64
4.2.1	La matrice di pixel	65
4.3	Alimentazione del chip	65
4.4	Elettronica di front-end di un generico chip di lettura	67
4.4.1	Componente analogica	67
4.4.2	Componente digitale	69
4.5	Front-end analogico di RD53A	71
4.5.1	Front-end lineare	72
4.5.2	Front-end differenziale	73
4.5.3	Front-end sincro	74
4.6	Front-end digitale di RD53A	77
4.6.1	Distributed Buffer Architecture	77
4.6.2	Central Buffer Architecture	77
4.7	Periferia del chip	79
4.7.1	Analog Chip Bottom	79
4.7.2	Digital Chip Bottom	80
4.8	Risultati e sviluppi futuri	81
5	Procedura di calibrazione del chip RD53A	83
5.1	Setup sperimentale	83
5.2	Operazioni preliminari	84
5.3	Il sistema di lettura YARR	86
5.4	Logica delle scansioni	88
5.4.1	Cicli eseguiti nelle scansioni	88
5.5	Procedura di configurazione del chip	92
5.5.1	Scan digitale	93
5.5.2	Scan analogico	93
5.5.3	Tuning della soglia globale	95
5.5.4	Tuning della soglia locale	96
5.5.5	Tuning del CSA	97
5.5.6	Tuning fine della soglia locale	97
5.5.7	Scan della soglia	100
5.5.8	Scan del ToT	101

6	Misure di caratterizzazione del chip RD53A	103
6.1	Linearità dei DAC	103
6.1.1	Isteresi della funzione di trasferimento	105
6.2	Impatto del tuning locale della soglia	106
6.3	Influenza della corrente di scarica sul ToT	108
6.4	Linearità del ToT	108
6.5	Scan del rumore	109
6.6	Curva I-V	111
6.7	Integrità dei bump-bond	112
6.8	Esposizione ad una sorgente radioattiva	114
Appendice A	Vincoli su λ dalle correzioni elettrodeboli al NLO al singolo Higgs	120
Appendice B	Il rivelatore ATLAS	125
Appendice C	Rivelatori a pixel monolitici	127
Conclusioni		129
Bibliografia		132

Introduzione

La possibilità di rivelare l'interazione di particelle con la materia è ampiamente sfruttata in numerose applicazioni, nella ricerca fondamentale in ambito accademico ma anche nell'industria. I rivelatori a pixel, appartenenti alla categoria dei dispositivi a semiconduttore, consentono di ottenere delle ottime prestazioni nella ricostruzione delle traiettorie descritte da particelle ionizzanti prodotte in esperimenti di fisica delle particelle di alta energia. La versatilità del loro utilizzo rappresenta uno dei fattori che hanno determinato il loro rapido sviluppo.

Nei rivelatori a semiconduttore il passaggio di una particella ionizzante, che determina una cessione di energia al mezzo con la conseguente produzione di un segnale elettrico, può essere localizzato con alta precisione grazie alla capacità di realizzare elettrodi estremamente granulari con le tecniche della microelettronica. L'idea di utilizzare tali dispositivi è nata intorno alla metà del XX secolo come alternativa alle camere a bolle, le quali risultavano inutilizzabili per la rivelazione di particelle ad alta frequenza di produzione. Tuttavia, la loro realizzazione era ostacolata da difficoltà di tipo tecnologico. L'interesse verso i rivelatori a semiconduttore ebbe un nuovo impulso intorno agli anni '70 del secolo scorso come conseguenza di due fattori: la nascita dei primi collisori adronici, caratterizzati da un elevato numero di particelle risultanti dall'interazione tra i fasci, e la scoperta del quark c , di vita media breve (dell'ordine del ps). I grandi progressi realizzati nel campo della fisica del Silicio, un particolare tipo di semiconduttore, hanno reso possibile la realizzazione di rivelatori a strip in grado di lavorare in ambienti caratterizzati da un elevato numero di particelle, anche nel caso in cui queste decadano dopo aver percorso distanze molto piccole, ad esempio dell'ordine di $100\text{ }\mu\text{m}$ nel caso di adroni contenenti i quark c o b . L'intensa attività di ricerca e sviluppo effettuata sul Silicio ed i progressi ottenuti nell'ambito della microelettronica hanno consentito, alcuni anni dopo, la realizzazione dei rivelatori a pixel. Essi consentono di ottenere risoluzioni spaziali dell'ordine di $10\text{ }\mu\text{m}$ sulla misura contemporanea delle due coordinate di un punto nel piano di misura. Grazie a queste caratteristiche uniche di precisione e granularità tali dispositivi sono tra quelli attualmente più adoperati per il tracciamento delle particelle in prossimità della regione di interazione tra i fasci accelerati ai collisori. I rivelatori a pixel ad oggi più diffusi sono quelli nei quali l'elettronica che provvede all'amplificazione ed al processamento del segnale rilasciato dalle particelle ionizzanti è integrata su un chip che è separato dalla regione nella quale avviene la generazione del segnale e successivamente collegato ad essa. Tali rivelatori sono utilizzati dall'esperimento ATLAS ad LHC e le loro prestazioni rendono possibili analisi di elevata sensibilità; tra queste, un esempio è rappresentato dalla ricerca di coppie di bosoni di Higgs.

L'esistenza di un campo scalare che determina la massa di fermioni e bosoni massivi, prevista teoricamente negli anni '60 del secolo scorso, è stata confermata nel 2012 grazie all'osservazione del bosone di Higgs, effettuata in maniera congiunta dagli esperimenti ATLAS e CMS, anch'esso ad LHC. Le costanti di accoppiamento di questa particella a fermioni e bosoni sono predette dal Modello Standard, ovvero dalla teoria che descrive l'interazione forte e quella elettrodebole;

la compatibilità di tali predizioni con i valori misurati ha contribuito a rafforzare la validità della teoria. Ciò nonostante, esiste un parametro del Modello Standard che non è stato ancora misurato: l'accoppiamento trilineare del campo di Higgs, λ . D'altra parte, diversi interrogativi non trovano risposta in questa teoria e richiedono una sua estensione. Tra i numerosi scenari che sono stati formulati nel corso degli anni, molti introducono nuovi fenomeni nel settore del bosone di Higgs, che rappresenta ancora la particella meno nota e dalle proprietà più singolari nel panorama che gli esperimenti agli acceleratori ci hanno permesso di conoscere. Un processo attraverso il quale si potrebbero manifestare effetti di nuova fisica, se esistenti, nel settore scalare e che è allo stesso tempo sensibile al parametro λ è rappresentato dalla produzione di coppie di bosoni di Higgs. Nel Modello Standard la sezione d'urto di produzione di due Higgs è talmente bassa da non essere accessibile alla luminosità alla quale LHC ha lavorato fino ad ora; tuttavia, l'esistenza di fisica non nota potrebbe incrementare la probabilità di osservare tale fenomeno. Per questo motivo la collaborazione ATLAS ha già svolto un'intensa attività di ricerca finalizzata alla rivelazione di tali eventi ed alla determinazione del coefficiente di accoppiamento trilineare dell'Higgs con l'obiettivo di verificare la compatibilità della sua misura con il valore predetto nell'ambito del Modello Standard.

In questa ricerca gioca un ruolo fondamentale l'identificazione dei jet prodotti da quark b , dal momento che essi sono presenti negli stati finali più abbondanti tra quelli originati da coppie di bosoni di Higgs. È quindi fondamentale riuscire a distinguere tali jet tra quelli di diversa natura che sono prodotti in grande quantità nelle interazioni tra i fasci di protoni accelerati ad LHC. L'efficienza di identificazione dei quark b ad ATLAS dipende in modo stringente dalle prestazioni del rivelatore di tracciamento, deputato alla misura di direzione, momento e carica delle particelle non neutre prodotte.

Il programma di misura e ricerche sulla produzione di coppie di Higgs sarà completato nel periodo di presa dati ad alta luminosità di LHC, durante il quale l'acceleratore fornirà in 12 anni di attività una quantità di collisioni protone - protone corrispondente ad una luminosità integrata pari a 4000 fb^{-1} , dieci volte superiore a quella accumulata dalla partenza di LHC fino al prossimo periodo di presa dati che si concluderà nel 2024. L'aumento di luminosità avverrà a partire dalla metà del 2027 e richiede la sostituzione del rivelatore di tracciamento attualmente utilizzato da ATLAS con un nuovo rivelatore, detto *Inner Tracker* (ITk). Esso sarà realizzato completamente in Silicio e conterrà un sottorivelatore a strip ed uno a pixel. L'Inner Tracker è stato progettato in modo tale che possa permettere l'identificazione dei quark b con prestazioni analoghe a quelle raggiunte fino ad ora ma in un ambiente più ostile, caratterizzato da una maggiore produzione di particelle a partire dalla collisione protone - protone come conseguenza dell'aumento di luminosità. La costruzione di questo rivelatore prende avvio oggi, a conclusione di un lavoro estremamente complesso di progettazione delle varie componenti e coinvolge un'ampissima comunità scientifica alla quale partecipano numerosi centri di ricerca e università che fanno parte della collaborazione ATLAS.

Il presente lavoro di tesi è incentrato sulla caratterizzazione di un prototipo dei chip di lettura che saranno utilizzati nel rivelatore a pixel dell'Inner Tracker. Per inquadrare le motivazioni che hanno spinto la collaborazione ATLAS a progettare l'ITk, nel capitolo 1 sono presentate le caratteristiche principali del bosone di Higgs e le modalità con cui coppie di queste particelle possono essere prodotte ed osservate ad LHC; inoltre, sono illustrati i risultati ottenuti fino ad ora da ATLAS nella determinazione del coefficiente di accoppiamento trilineare del bosone di Higgs tramite la ricerca diretta di coppie di tali particelle. Il capitolo 2 è dedicato al problema dell'identificazione dei quark b ad ATLAS: sono descritti gli algoritmi adoperati, l'impatto avuto

dalle modifiche apportate al rivelatore di tracciamento attualmente installato e quello previsto per l'Inner Tracker. Inoltre, è illustrata l'estrapolazione dei risultati della ricerca di coppie di Higgs descritta nel capitolo 1 ottenuta considerando l'aumento di luminosità e la nuova geometria del rivelatore di tracciamento che caratterizzeranno ATLAS a partire dal 2027. Nel capitolo 3 sono descritte le proprietà generali dei rivelatori a semiconduttore e di quelli a pixel di Silicio; l'ultima sezione è dedicata al confronto tra le caratteristiche dei rivelatori a pixel attualmente utilizzati da ATLAS e di quelli previsti per ITk. Il capitolo 4 illustra la struttura e la logica di processamento dei dati del chip RD53A, il prototipo dei chip di lettura del futuro rivelatore di tracciamento analizzato in questa tesi. Le misure di calibrazione effettuate sono oggetto del capitolo 5, nel quale è descritto anche il setup sperimentale utilizzato. Esso è stato allestito nel Laboratorio INFN "Costruzione grandi apparati" del Dipartimento di Matematica e Fisica "E. De Giorgi" presso l'Università del Salento. La sezione INFN di Lecce, infatti, è impegnata nell'assemblaggio di alcune componenti dei rivelatori a pixel che saranno installati nell'Inner Tracker. La loro produzione sarà accompagnata da test di funzionalità e da campagne di misure volte a certificare la loro qualità prima dell'integrazione nel rivelatore. Le misure condotte sul chip ed illustrate in questa tesi rappresentano un primo esercizio di implementazione di un sistema di acquisizione dati per i rivelatori a pixel di ATLAS fondamentale nella messa a punto del setup sperimentale che sarà utilizzato durante la fase di produzione. Il capitolo 6 descrive delle misure originali effettuate sul chip RD53A che permettono di valutare in dettaglio la sua risposta; esso comprende, inoltre, i risultati dell'analisi della radiazione emessa da una sorgente di particelle β .

Capitolo 1

La produzione di coppie di Higgs

Uno degli obiettivi dei programmi di ricerca degli esperimenti ATLAS (*A Toroidal LHC Apparatus*, [1]) e CMS (*Compact Muon Solenoid*, [2]) ad LHC (*Large Hadron Collider*, [3]) è rappresentato dalla determinazione del coefficiente di accoppiamento trilineare del bosone di Higgs H . Questo può essere misurato tramite l'osservazione di coppie HH ; la loro produzione, tuttavia, è associata ad una sezione d'urto talmente bassa da non consentire la loro osservazione.

Effetti di fisica oltre il Modello Standard, quali la produzione attraverso il decadimento di nuove particelle con elevata sezione d'urto o valori anomali degli accoppiamenti trilineari e quadrilineari, potrebbero determinare una frequenza di produzione di coppie di Higgs osservabile. Per questo motivo i dati raccolti fino ad ora dalle due collaborazioni sono stati analizzati con l'obiettivo di ricercare una produzione anomala di coppie HH .

In questo capitolo sono introdotte le principali caratteristiche del bosone di Higgs ed i vari canali attraverso i quali le coppie HH possono essere prodotte; inoltre, sono descritte le ricerche di produzione anomala svolte da ATLAS ed i relativi risultati.

1.1 Il bosone di Higgs

Nel 1964 Brout ed Englert [4] ed indipendentemente Higgs [5] proposero una teoria che consentiva di unificare l'interazione debole e quella elettromagnetica e di spiegare l'acquisizione della massa da parte delle particelle. Tale teoria, nota come *meccanismo BEH* (dalle iniziali dei fisici che l'hanno introdotta), prevede che la fisica accessibile ai moderni acceleratori rappresenti il residuo della rottura spontanea della simmetria $SU(3)_C \times SU(2)_W \times U(1)_Y$, che descrive il Modello Standard, nella componente $SU(2)_W \times U(1)_Y$. Il principio alla base del meccanismo BEH consiste nel postulare l'esistenza di quattro campi scalari reali organizzati come un doppietto di campi complessi nello spazio $SU(2)$, detto *doppietto di Higgs*, e come un singoletto nello spazio di colore e di ipercarica. Il doppietto di Higgs Φ , di ipercarica $Y = 1$, può essere scritto come

$$\Phi = \begin{pmatrix} \phi^+ \\ \phi^0 \end{pmatrix}$$

ed il potenziale V ad esso associato assume l'espressione

$$V(\Phi^\dagger \Phi) = \mu^2 \Phi^\dagger \Phi + \lambda (\Phi^\dagger \Phi)^2$$

in cui μ e λ rappresentano due costanti di accoppiamento. Tale potenziale rispetta la simmetria di gauge della Lagrangiana del Modello Standard ma ha un minimo degenere. La realizzazione di uno degli infiniti stati di minimo tra loro equivalenti determina la rottura spontanea della simmetria esatta $SU(2)_W \times U(1)_Y$ senza agire su $SU(3)_C$.

Un possibile stato di vuoto è rappresentato da

$$\frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ v \end{pmatrix}$$

in cui solo la parte reale della componente neutra del doppietto di Higgs assume un valore diverso da 0. Tale stato risulta invariante sotto un nuovo gruppo di simmetria $U(1)$ associato all'operatore di carica $Q = T_3 + Y/2$, in cui T_3 rappresenta l'autovalore di isospin debole. La simmetria associata all'operatore Q è quella elettromagnetica; pertanto, il gruppo di simmetria che descrive il Modello Standard in seguito alla realizzazione di uno dei possibili stati di minimo è $SU(3)_C \times U(1)_{EM}$.

Espandendo il potenziale intorno al minimo scelto si ottiene l'espressione

$$V(H) = \frac{1}{2}m_H^2 H^2 + \lambda v H^3 + \frac{1}{4}\lambda H^4 \quad (1.1)$$

in cui H è l'unico stato fisico associato al doppietto di Higgs che sopravvive alla rottura e prende pertanto il nome di *campo di Higgs*, di massa m_H .

I termini di massa per i bosoni di gauge $W^\pm$ e Z sono generati naturalmente, nello sviluppo della Lagrangiana intorno allo stato di minima energia, dalle derivate covarianti del campo di Higgs, le quali garantiscono l'invarianza di gauge del termine cinetico della Lagrangiana del campo scalare. In questo stesso processo i termini contenenti gli altri tre campi del doppietto di Higgs, relativi ai valori di aspettazione nulli, si cancellano coerentemente con il fatto che non rappresentano gradi di libertà effettivi ma solo libertà di gauge. La scelta dei numeri quantici dello stato di vuoto prescelto e quindi del campo H garantisce che la simmetria residua $U(1)_{EM}$ rimanga non rotta e pertanto il fotone (la combinazione dei campi di gauge associata a tale simmetria) non acquista massa. Infine, i termini di massa per i fermioni sono prodotti da interazioni di Yukawa del campo di Higgs proporzionali alla combinazione $\bar{f}_L f_R + \bar{f}_R f_L$ (in cui L e R rappresentano le componenti di chiralità sinistrorsa, *left*, e destrorsa, *right*, del fermione f) invarianti sotto il gruppo di simmetria $SU(2)_W$. Il meccanismo BEH permette così l'introduzione di tutte le masse conservando la simmetria di gauge della Lagrangiana che garantisce la rinormalizzabilità del Modello Standard.

Nei decenni successivi alla formulazione della teoria diversi esperimenti hanno cercato di individuare la particella associata al campo teorizzato da Brout, Englert ed Higgs per poter confermare la validità della loro ipotesi di rottura della simmetria elettrodebole. Nel 2012 le collaborazioni ATLAS e CMS hanno annunciato l'osservazione di una particella compatibile con quella prevista dal meccanismo BEH ([6, 7]); da allora sono stati condotti diversi studi volti a determinare le proprietà di tale particella, per stabilire se si trattasse effettivamente di quella teorizzata.

La massa del bosone di Higgs non è predetta dal meccanismo BEH; su di essa sono stati posti dei vincoli che derivano dai risultati delle misure elettrodeboli, che ne dipendono logicamente, e della richiesta di consistenza della teoria. La massa della particella scoperta nel

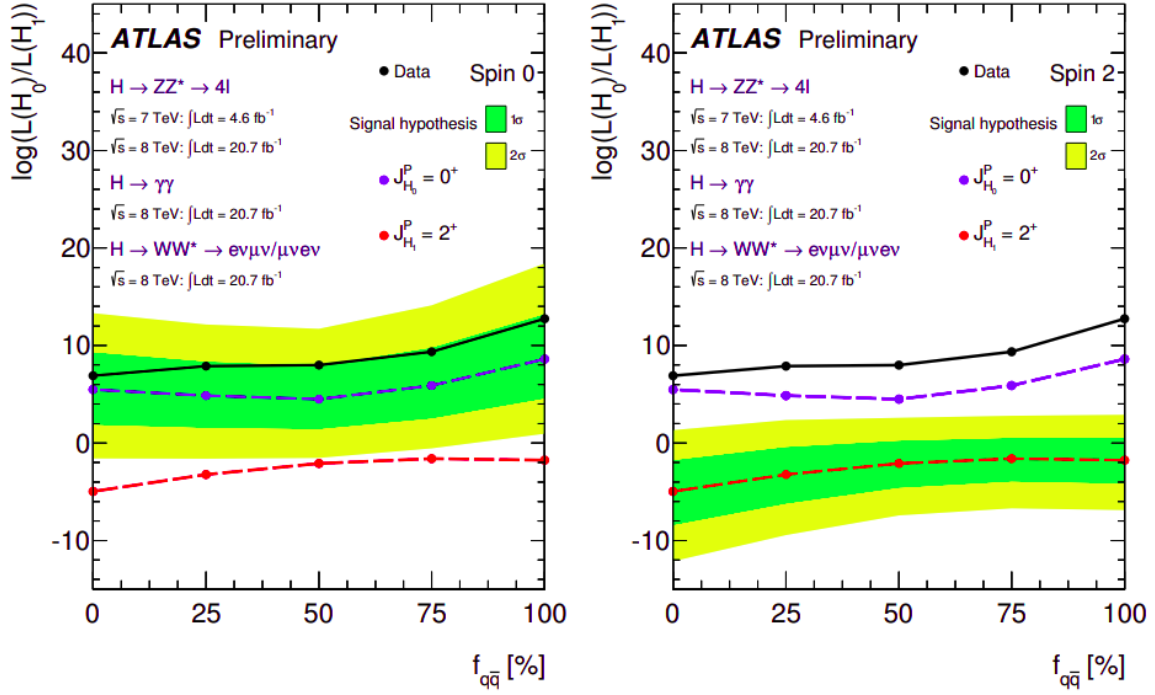


Figura 1.1: Fit di massima verosimiglianza ai dati nell'ipotesi che la particella individuata da ATLAS e CMS nel 2012 abbia spin 0 (sinistra) oppure 2 (destra). Da [9].

2012, pari a (125.10 ± 0.14) GeV [8], risulta compatibile con tali limiti. Lo spin del campo di Higgs è invece una predizione fondamentale; pertanto, la determinazione dello spin della particella individuata può essere determinante nell'escludere che essa coincida con quella teorizzata. In figura 1.1 è riportato il risultato di uno studio dello spin condotto da ATLAS [9]. I due grafici rappresentano il risultato di test di massima verosimiglianza ai dati delle ipotesi di spin 0 (a sinistra) e di spin 2⁽¹⁾ (a destra) in funzione di $f_{q\bar{q}}$; tale parametro tiene conto del fatto che la produzione di una particella di spin 2 può avvenire mediante l'annichilazione di due fermioni (in alternativa all'interazione di due gluoni) con una frequenza $f_{q\bar{q}}$ che dipende dal modello specifico e determina la probabilità complessiva che gli eventi osservati siano compatibili con l'ipotesi di una particella di spin 2. La linea nera rappresenta il valore della *test statistic*, definita come logaritmo del rapporto delle likelihood delle osservabili sperimentali (distribuzioni di variabili angolari dei prodotti di decadimento) secondo le due ipotesi valutato per i dati, quella blu l'attesa nell'ipotesi di un segnale di spin 0 e quella rossa l'attesa nell'ipotesi di un segnale di spin 2; la banda verde e quella gialla rappresentano gli intervalli di confidenza relativi all'ipotesi considerata rispettivamente ad 1σ e 2σ . Nel grafico a sinistra i dati raccolti sono compatibili con l'ipotesi di spin 0 entro 1σ indipendentemente dal valore di $f_{q\bar{q}}$, mentre nel secondo caso sono sempre all'esterno dell'intervallo di confidenza a 2σ ; pertanto, l'ipotesi di spin 2 per la nuova particella è rigettata con un livello di confidenza superiore al 99.9%. Analogamente a quelle sullo spin, le misure delle costanti di accoppiamento della particella a fermioni e bosoni hanno fornito risultati compatibili con quelli attesi nel Modello Standard. Le ricerche effettuate dalle due collaborazioni hanno pertanto consentito di concludere che la particella individuata è il bosone che permette di convalidare la teoria di rottura della simmetria elettrodebole.

⁽¹⁾Le ipotesi di spin considerate sono 0 e 2 in quanto tale analisi è stata condotta osservando il decadimento dell'Higgs in coppie di bosoni di spin 1 ($\gamma\gamma$, W^+W^- , ZZ).

1.1.1 Canali di produzione e di decadimento

In figura 1.2 sono rappresentati i principali canali di produzione dell'Higgs. Il grafico in alto a sinistra è relativo al processo di *gluon-gluon fusion* (ggF, $pp \rightarrow H$), quello in alto a destra al processo di *vector boson fusion* (VBF, $pp \rightarrow qq'H$, in cui V può rappresentare il bosone W o quello Z), quello in basso a sinistra è il processo di *Higgs-strahlung* (VH, $pp \rightarrow VH$) mentre l'ultimo riguarda la produzione dell'Higgs in associazione ai quark $t\bar{t}$ ($t\bar{t}H$, $pp \rightarrow t\bar{t}H$). Nei

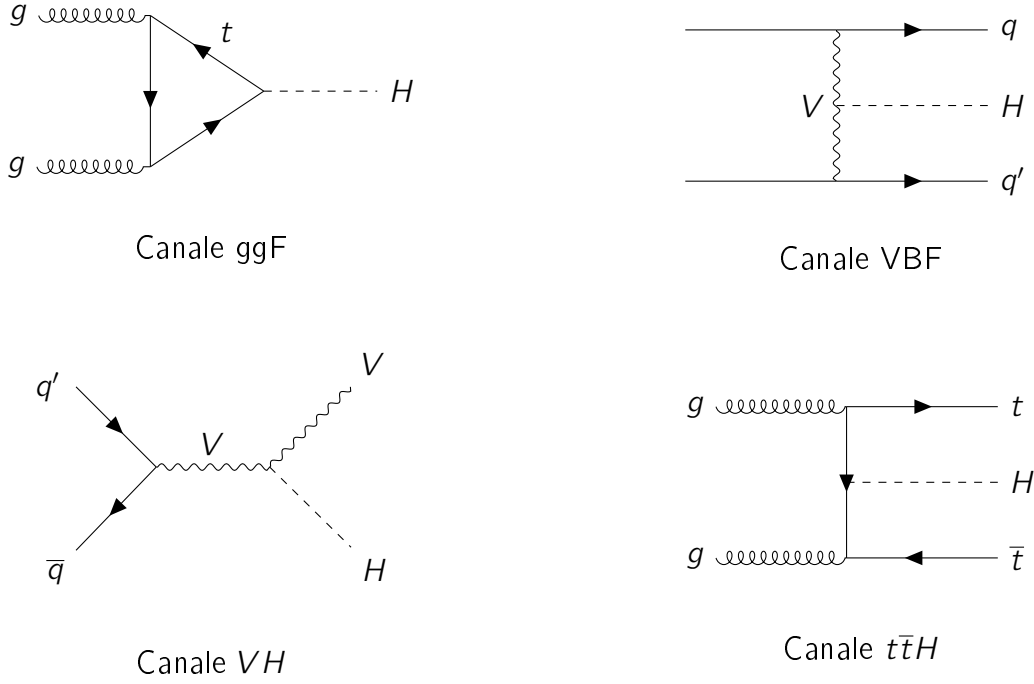


Figura 1.2: Esempi di processi di produzione del bosone di Higgs.

canali che prevedono l'accoppiamento dell'Higgs ad un fermione si sono considerati quelli che coinvolgono il quark t perché più probabili, in quanto la costante di accoppiamento ai fermioni è proporzionale alla loro massa.

In figura 1.3 (sinistra) sono rappresentate le sezioni d'urto predette per tali processi in funzione dell'energia del centro di massa nell'interazione protone - protone (pp), $\sqrt{s}$, con l'indicazione del livello di precisione del calcolo perturbativo. In tabella 1.1 sono estrapolati i valori relativi a $\sqrt{s} = 7$ TeV, 8 TeV, 13 TeV e 14 TeV; le prime due rappresentano le energie alle quali LHC ha lavorato nel periodo di presa dati indicato come *run 1*, svoltosi tra il 2011 ed il 2012, la terza è l'energia corrispondente al *run 2*, svoltosi dal 2015 al 2018, e l'ultima è l'energia prevista per il *run 3*, programmato dal 2022 al 2024. Risulta evidente come il canale di produzione predominante sia quello della gluon-gluon fusion.

La larghezza del bosone di Higgs nel Modello Standard è pari a $\Gamma_H = 4.07 \times 10^{-3}$ GeV, calcolata assumendo una massa di 125 GeV [10]. In figura 1.3 (destra) è rappresentato il rapporto tra la larghezza di decadimento in ciascuno dei possibili canali rispetto a quella totale (indicata come *branching ratio*) in funzione della massa dell'Higgs ed in tabella 1.2 sono indicati i valori estratti in corrispondenza di $m_H = 125$ GeV. Il canale di decadimento più probabile è quello in $b\bar{b}$.

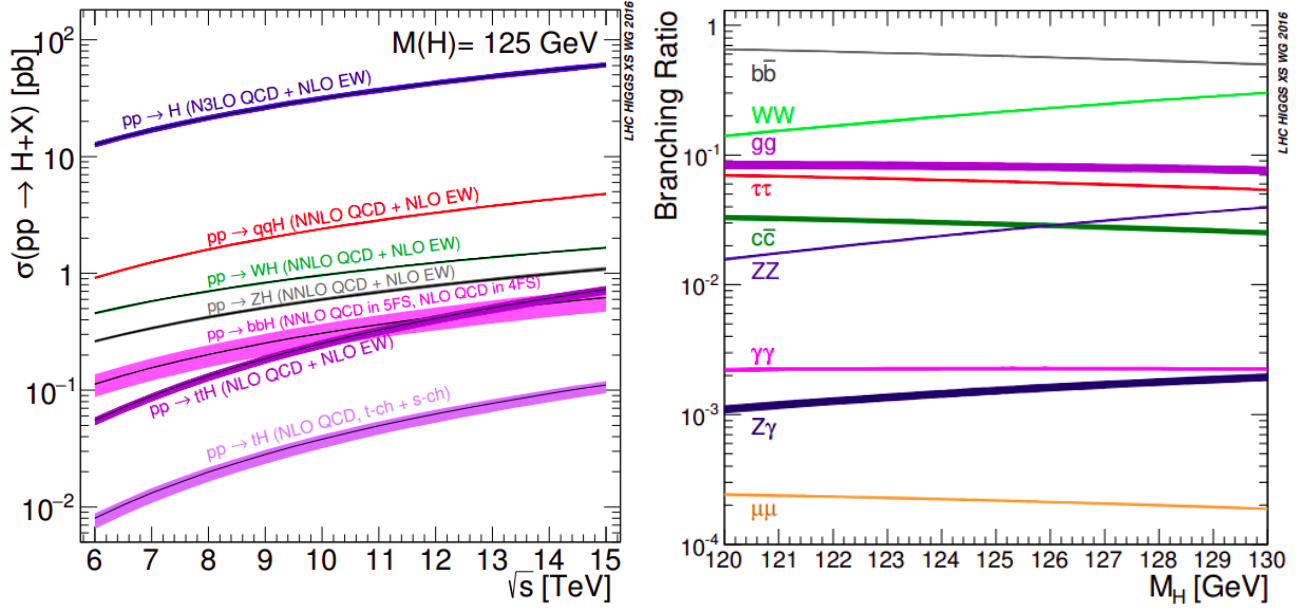


Figura 1.3: Sinistra: sezione d'urto dei vari canali di produzione dell'Higgs in funzione dell'energia nel centro di massa. Destra: branching ratio dei principali canali di decadimento dell'Higgs in funzione della sua massa. Da [8].

$\sqrt{s}$ (TeV)	Sezione d'urto di produzione [pb]					
	ggF	VBF	WH	ZH	$t\bar{t}H$	totale
7	$16.9^{+4.4\%}_{-7.0\%}$	$1.24^{+2.1\%}_{-2.1\%}$	$0.58^{+2.2\%}_{-2.3\%}$	$0.34^{+3.1\%}_{-3.0\%}$	$0.09^{+5.6\%}_{-10.2\%}$	19.1
8	$21.4^{+4.4\%}_{-6.9\%}$	$1.60^{+2.3\%}_{-2.1\%}$	$0.70^{+2.1\%}_{-2.2\%}$	$0.42^{+3.4\%}_{-2.9\%}$	$0.13^{+5.9\%}_{-10.1\%}$	24.2
13	$48.6^{+4.6\%}_{-6.7\%}$	$3.78^{+2.2\%}_{-2.2\%}$	$1.37^{+2.6\%}_{-2.6\%}$	$0.88^{+4.1\%}_{-3.5\%}$	$0.50^{+6.8\%}_{-9.9\%}$	55.1
14	$54.7^{+4.6\%}_{-6.7\%}$	$4.28^{+2.2\%}_{-2.2\%}$	$1.51^{+1.9\%}_{-2.0\%}$	$0.99^{+4.1\%}_{-3.7\%}$	$0.60^{+6.9\%}_{-9.8\%}$	62.1

Tabella 1.1: Sezione d'urto dei principali processi di produzione dell'Higgs alle energie alle quali ha lavorato LHC durante i run 1 ($\sqrt{s} = 7$ TeV e $\sqrt{s} = 8$ TeV) e 2 ($\sqrt{s} = 13$ TeV) ed all'energia prevista per il run 3 ($\sqrt{s} = 14$ TeV). Da [8].

Canale di decadimento	Branching ratio	Incertezza relativa
$H \rightarrow b\bar{b}$	5.82×10^{-1}	$+1.2\%$ -1.3%
$H \rightarrow W^+W^-$	2.14×10^{-1}	$\pm 1.5\%$
$H \rightarrow \tau^+\tau^-$	6.27×10^{-2}	$\pm 1.6\%$
$H \rightarrow c\bar{c}$	2.89×10^{-2}	$+5.5\%$ -2.0%
$H \rightarrow ZZ$	2.62×10^{-2}	$\pm 1.5\%$
$H \rightarrow \gamma\gamma$	2.27×10^{-3}	2.1%
$H \rightarrow Z\gamma$	1.53×10^{-3}	$\pm 5.8\%$
$H \rightarrow \mu^+\mu^-$	2.18×10^{-4}	$\pm 1.7\%$

Tabella 1.2: Branching ratio dei principali processi di decadimento dell'Higgs con incertezza relativa in corrispondenza di $m_H = 125$ GeV. Da [8].

1.1.2 L'accoppiamento trilineare λ

Nonostante gli anni successivi alla scoperta del bosone di Higgs abbiano permesso di determinare con precisione sempre crescente molte delle sue proprietà, una quantità ancora non misurata è il coefficiente di accoppiamento trilineare dell'Higgs, detto anche di *self-coupling*, ovvero il termine λ nell'espressione (1.1). Tale misura è indispensabile affinché si possano porre dei vincoli alla forma del potenziale di Higgs nelle vicinanze del minimo; pertanto la sua determinazione è un obiettivo primario del programma futuro degli esperimenti ad LHC.

La misura diretta del coefficiente di proporzionalità di un termine di ordine n in H è legata all'osservazione di $n - 1$ bosoni di Higgs; pertanto, considerando che il fattore λ è legato alle potenze H^3 e H^4 (come indicato nella relazione (1.1)) è possibile ricavare informazioni su di esso da processi che coinvolgono la produzione di 2 o 3 Higgs. Come conseguenza dell'elevata massa dello stato finale, la sezione d'urto di produzione di 3 Higgs è così bassa (dell'ordine dei decimi di fb) che l'osservazione di tale processo è impensabile alla luminosità⁽²⁾ di LHC ed anche a quella di *High Luminosity-LHC* (HL-LHC), termine con il quale si indica il periodo di presa dati previsto a partire dal 2027, caratterizzato da un aumento della luminosità. Maggiori speranze sono invece riposte nella produzione di coppie di Higgs, sulla cui ricerca si stanno concentrando i due esperimenti ATLAS e CMS. La determinazione del coefficiente di self-coupling è complicata da due fattori: la bassa sezione d'urto dei processi che determinano la produzione di coppie di Higgs, ed il fatto che alcuni di essi hanno una dipendenza blanda o addirittura nulla da λ . Tali caratteristiche saranno illustrate con maggiore dettaglio nella sezione 1.2.

Lo studio dell'accoppiamento trilineare può essere svolto considerando una teoria di campo effettiva (*Effective Field Theory*, EFT) che estenda il Modello Standard. Tale procedura trova giustificazione nel fatto che fino ad ora non è stata osservata alcuna nuova fisica ad LHC: è possibile pensare che la presenza di nuove particelle si manifesti a scale ben superiori al TeV e parametrizzare dunque i loro effetti tramite un insieme di operatori di dimensione 6, che si suppone abbiano effetti predominanti rispetto a quelli di dimensione 5⁽³⁾. Tali operatori determinano il verificarsi di nuove interazioni ed una modifica delle costanti di accoppiamento relative ai processi già noti all'interno del Modello Standard. I dati raccolti dalle due collaborazioni possono dunque essere esaminati in termini di questi nuovi parametri. In generale in una EFT il numero di nuovi parametri è elevato; inoltre le analisi che ricercano la produzione di coppie di bosoni di Higgs hanno una sensibilità limitata a causa delle basse sezioni d'urto e degli abbondanti processi di fondo. Nell'interpretazione dei dati si ricorre pertanto all'identificazione di processi sensibili solo a pochi parametri, oppure ad ipotesi semplificative analizzando le misure in termini di uno o due parametri della Lagrangiana EFT alla volta mentre si fissano gli altri al valore che assumono nel Modello Standard. L'approccio dell'EFT alla ricerca di deviazioni dalle predizioni del Modello Standard ha il vantaggio di non essere dipendente da nessuno specifico modello di nuova fisica.

Naturalmente numerose analisi dei dati di LHC sono effettuate seguendo l'approccio del test di particolari ipotesi teoriche; ciò accade anche nello studio della produzione di coppie di bosoni di Higgs. Tra le varie teorie di fisica oltre il Modello Standard (*Beyond Standard Model*, BSM) formulate fino ad ora, due sono quelle più studiate in relazione alla produzione di coppie di Higgs:

⁽²⁾Si definisce *luminosità istantanea*, $\mathcal{L}$, il rapporto tra la frequenza di eventi, f , e la relativa sezione d'urto σ : $\mathcal{L} = \frac{f}{\sigma}$. La *luminosità integrata*, L , rappresenta l'integrale della precedente sul tempo di presa dati: $L = \frac{N}{\sigma}$, in cui N indica il numero di eventi verificatisi.

⁽³⁾La dimensione degli operatori della Lagrangiana del Modello Standard è 4.

quella del *two Higgs Doublet Model* (2HDM) e quella della *Gravity-mediated Dark Matter*. La 2HDM prevede l'esistenza di un doppietto di Higgs che si affianchi a quello già esistente nel Modello Standard. Degli 8 gradi di libertà presenti in tale settore di Higgs esteso durante la fase esatta del Modello Standard, in seguito alla rottura elettrodebole 3 vengono utilizzati per conferire massa ai bosoni W^+ , W^- e Z mentre i restanti 5 sono associati ad altrettanti stati fisici: due scalari (h ed H), uno pseudoscalare (A) e due stati carichi ($H^\pm$). Assumendo che la particella osservata nel 2012 corrisponda allo stato scalare di minore massa, è possibile ipotizzare la produzione risonante di coppie di tali particelle a partire dallo stato scalare di massa maggiore. La teoria della Gravity-mediated Dark Matter individua nel *gravitone*, una particella massiva di spin 2, il responsabile dell'interazione della materia oscura con l'Universo conosciuto; tale gravitone potrebbe determinare la produzione risonante di coppie di Higgs.

Un metodo per porre dei vincoli al valore di k_λ alternativo alla ricerca di coppie di bosoni di Higgs consiste nello studio delle correzioni elettrodeboli al NLO dei processi di produzione non risonante e di decadimento di un singolo Higgs, nei quali interviene il termine λ . Maggiori informazioni sono fornite nell'Appendice A.

1.2 Produzione di coppie di Higgs in collisioni pp

La produzione di coppie di bosoni di Higgs in collisioni pp , come quelle che si verificano ad LHC, può avvenire attraverso diversi canali; in figura 1.4 è rappresentato un possibile diagramma di Feynman associato a ciascuno di essi. I diversi canali di produzione sono:

- gluon-gluon fusion (in alto a sinistra);
- vector boson fusion (in alto a destra);
- produzione associata al singolo quark top ($tjHH$, caso particolare di VBF in cui uno dei due quark nello stato finale è di tipo t);
- produzione associata ad un bosone vettore V (VHH , in basso a sinistra);
- produzione associata a $t\bar{t}$ ($t\bar{t}HH$, in basso a destra).

Nell'ipotesi che l'accoppiamento trilineare dell'Higgs sia quello previsto dal Modello Standard sono state calcolate le predizioni al NLO (*Next-to-Leading Order*) per le sezioni d'urto nei singoli canali di produzione elencati in funzione dell'energia nel centro di massa, rappresentate in figura 1.5 (sinistra). In tabella 1.3 sono estratti i valori relativi alle energie $\sqrt{s} = 8$ TeV, 13 TeV e 14 TeV per i due canali di produzione predominanti, ggF e VBF; l'errore sulla predizione teorica è separato in due contributi per isolare (nel secondo) la componente dovuta alle incertezze di natura sperimentale sulle funzioni di distribuzione dei partoni (*Parton Distribution Functions*, PDF) nel protone. Il canale predominante è quello della gluon-gluon fusion come nel caso del singolo Higgs, ma la sezione d'urto è inferiore di 3 ordini di grandezza.

I valori relativi all'energia $\sqrt{s} = 14$ TeV sono stati calcolati al LO (*Leading Order*) ed al NLO in funzione del rapporto λ/λ^{SM} , in cui λ^{SM} rappresenta il valore del coefficiente di self-coupling nel Modello Standard, fissando al valore assunto nel Modello Standard gli accoppiamenti che risultano dall'utilizzo della EFT; i risultati ottenuti sono rappresentati in figura 1.5 (destra). Le linee tratteggiate e le bande chiare corrispondono alle previsioni al LO, le linee continue e

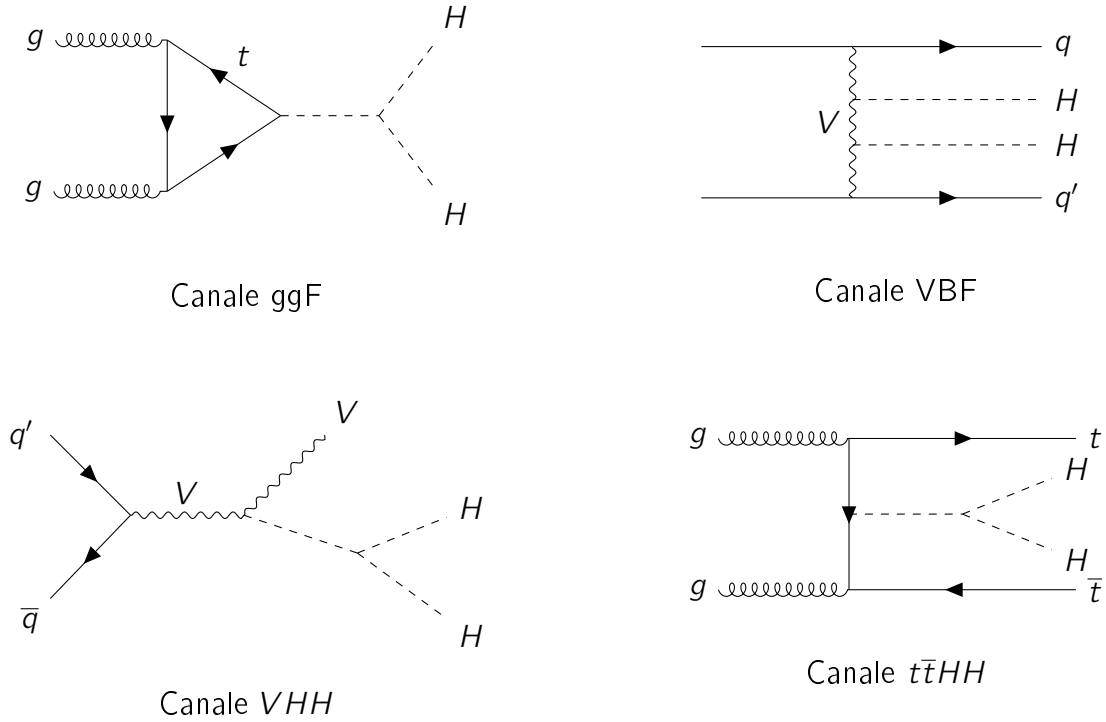


Figura 1.4: Esempi di processi di produzione di coppie di bosoni di Higgs dall'interazione pp.

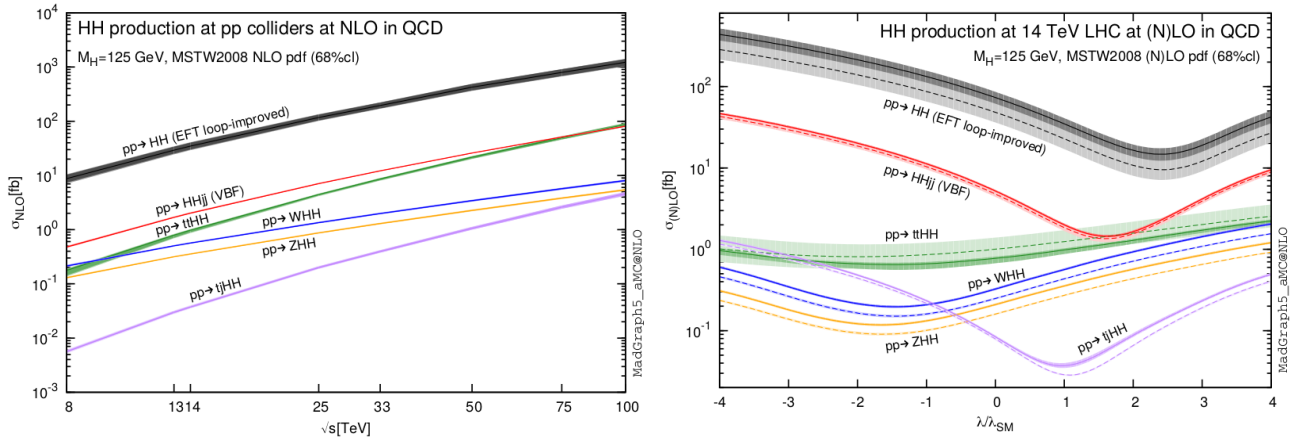


Figura 1.5: Sinistra: sezione d'urto al NLO dei principali canali di produzione di coppie di Higgs in funzione dell'energia nel centro di massa. Destra: sezione d'urto dei principali canali di produzione di coppie di Higgs in funzione del rapporto $\lambda/\lambda_{\text{SM}}$. Le linee tratteggiate e le bande chiare corrispondono alle previsioni al LO, le linee continue e le bande scure rappresentano le previsioni al NLO. Da [11].

le bande scure rappresentano quelle al NLO. Si osserva come gli errori sulle predizioni al NLO siano inferiori rispetto a quelli nei calcoli all'ordine perturbativo più basso.

Nelle prossime sezioni sono forniti maggiori dettagli riguardo ai due canali di produzione predominanti.

Sezione d'urto di produzione al NLO [fb]			
Canale	$\sqrt{s} = 8 \text{ TeV}$	$\sqrt{s} = 13 \text{ TeV}$	$\sqrt{s} = 14 \text{ TeV}$
ggF	$8.73^{+17+2.9\%}_{-16-3.7\%}$	$29.3^{+15+2.1\%}_{-14-2.5\%}$	$34.8^{+15+2.0\%}_{-14-2.5\%}$
VBF	$0.479^{+1.8+2.8\%}_{-1.8-2.0\%}$	$1.684^{+1.4+2.6\%}_{-0.9-1.9\%}$	$2.017^{+1.3+2.5\%}_{-1.0-1.9\%}$

Tabella 1.3: Sezioni d'urto al NLO (in fb) dei principali canali di produzione di coppie di Higgs. L'errore sulla predizione teorica è separato in due contributi per isolare (nel secondo) la componente dovuta alle incertezze di natura sperimentale sulle PDF dei partoni nel protone. Da [11].

1.2.1 Gluon-gluon fusion

I due diagrammi che contribuiscono alla sezione d'urto nel canale gluon-gluon fusion nel Modello Standard sono rappresentati in figura 1.6. Il basso valore della sezione d'urto associata a

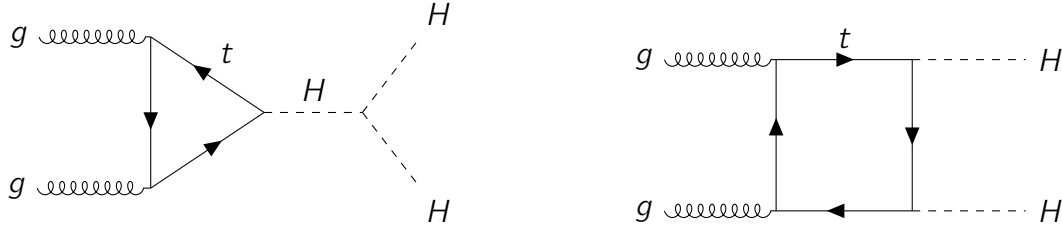


Figura 1.6: Diagrammi che contribuiscono alla produzione di coppie di bosoni di Higgs nel canale gluon-gluon fusion nel Modello Standard.

questo canale è determinato dal fatto che tali diagrammi interferiscono negativamente; inoltre, si osserva come solo nel primo grafico intervenga l'accoppiamento trilineare, e ciò indebolisce la dipendenza della sezione d'urto da tale quantità. Tali fattori causano la difficoltà della determinazione di λ alla quale si è accennato nella sezione 1.1.2.

In EFT, i nuovi operatori che estendono la Lagrangiana fanno sì che la produzione di coppie in gluon-gluon fusion avvenga tramite i diagrammi in figura 1.7. I primi due diagrammi sono analoghi a quelli presenti nel Modello Standard ma sono descritti da diversi valori delle costanti λ e y_t , quest'ultima relativa all'accoppiamento dell'Higgs con il quark t ; è possibile dunque introdurre i coefficienti $k_\lambda = \lambda/\lambda^{SM}$ e $k_t = y_t/y_t^{SM}$, che quantificano la deviazione delle nuovi costanti da quelle del Modello Standard, indicate con l'apice SM . I tre diagrammi nella parte inferiore della figura sono assenti nel Modello Standard e descrivono l'interazione tra due gluoni e due bosoni di Higgs, descritta dalla costante c_{2g} , tra due gluoni ed un bosone di Higgs, descritta dalla costante c_g , e tra due quark e due bosoni di Higgs, descritta dalla costante c_2 . Come già accennato, i risultati in figura 1.5 sono stati ottenuti fissando tutti i coefficienti ad eccezione di k_λ al valore che assumono all'interno del Modello Standard.

1.2.2 Vector boson fusion

Il canale di produzione vector boson fusion ha una sezione d'urto inferiore rispetto a quella della gluon-gluon fusion ed inoltre l'accoppiamento trilineare ha un impatto minore su di esso; tuttavia, offre la possibilità di investigare gli accoppiamenti VVH e $VVHH$, come rappresentato in figura 1.8, nella quale i primi 3 diagrammi sono relativi alla EFT.

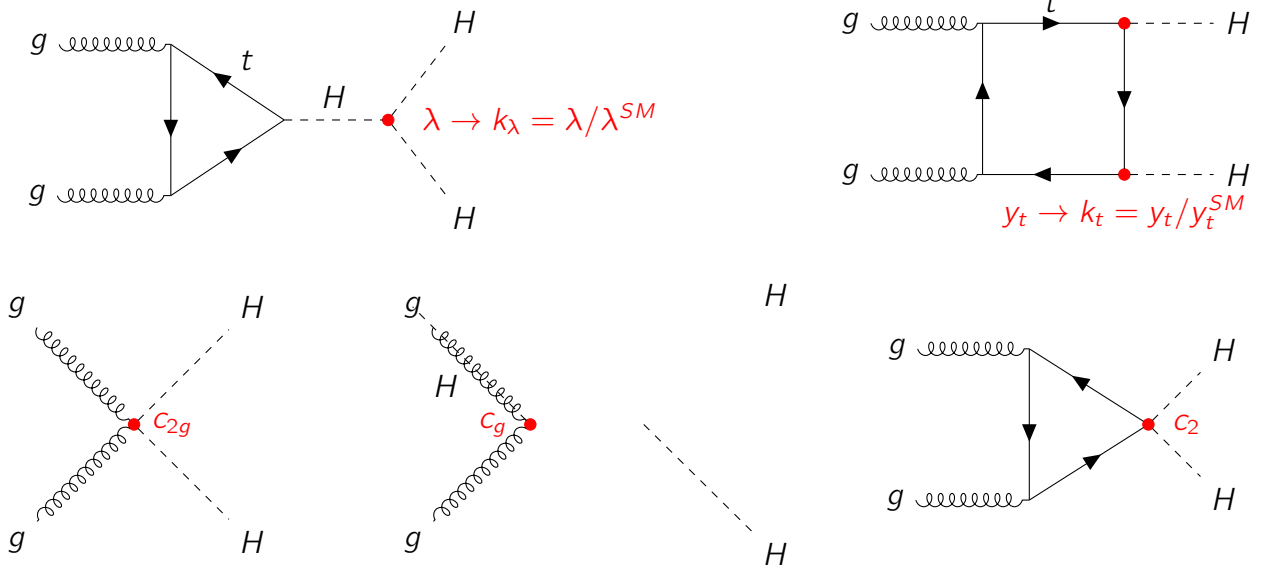


Figura 1.7: Processi che contribuiscono alla produzione di coppie di bosoni di Higgs nel canale gluon-gluon fusion nell'ambito della teoria di campo effettiva.

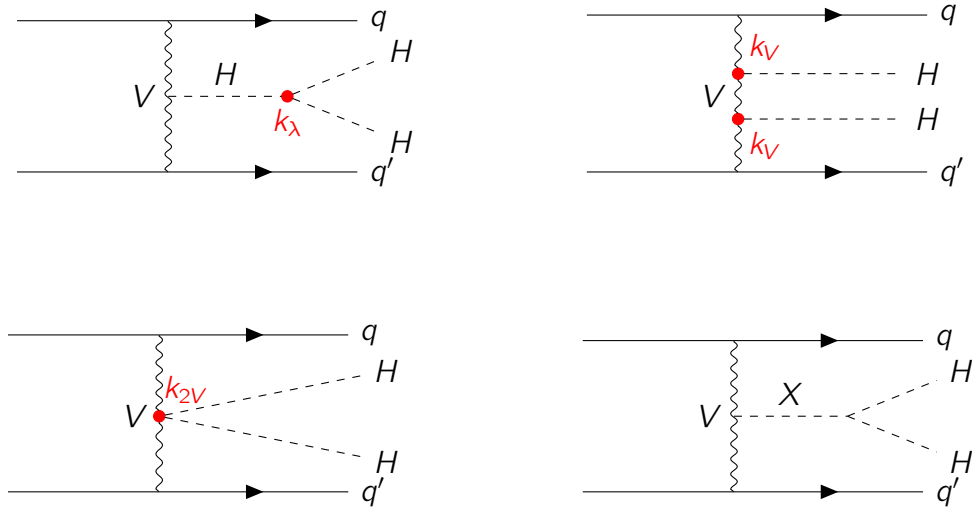


Figura 1.8: Processi che contribuiscono alla produzione di coppie di bosoni di Higgs nel canale vector boson fusion nell'ambito della teoria di campo effettiva (primi tre diagrammi) ed in teorie BSM (ultimo diagramma).

Pertanto, oltre al fattore k_λ già introdotto, in tale canale intervengono anche k_V e k_{2V} ; quest'ultimo in particolare descrive una nuova interazione efficace che può essere testata solo nel canale VBF. Anche tali accoppiamenti sono fissati al valore che assumono all'interno del Modello Standard per ottenere i risultati in figura 1.5. L'ultimo diagramma in figura 1.8 emerge da una possibile teoria BSM che prevede la produzione risonante di coppie di bosoni di Higgs in canale VBF a partire dalla particella esotica⁽⁴⁾ X .

⁽⁴⁾Non appartenente al Modello Standard.

1.3 La ricerca di coppie di Higgs ad LHC

La ricerca di coppie di Higgs ad LHC si è concentrata fino ad oggi sul canale di produzione predominante, ovvero la gluon-gluon fusion; i dati sono stati analizzati con procedure separatamente ottimizzate per la ricerca dei diversi possibili decadimenti dell'Higgs. Nell'ipotesi che non esista fisica oltre il Modello Standard la quantità di dati raccolti sino ad ora non consente l'osservazione di coppie di bosoni di Higgs, dato che la sezione d'urto di questo processo è dell'ordine di poche decine di fb e le efficienze di selezione sono troppo basse, a causa della necessità di abbattere gli abbondanti processi di fondo. La presenza di fisica BSM potrebbe invece aumentare la frequenza di produzione di coppie di Higgs, rendendola accessibile anche alla luminosità di LHC; pertanto, quello che si ricerca in queste analisi è una produzione anomala di coppie HH , parametrizzata dalla EFT nel caso di produzione non risonante o da uno specifico modello nel caso di produzione risonante.

Le analisi condotte fino ad ora hanno ottenuto limiti superiori alla sezione d'urto di produzione di coppie di Higgs; i canali più sensibili sono stati successivamente sfruttati per vincolare il valore di k_λ . Nel seguito di questo paragrafo si descriveranno le analisi più importanti e si accennerà ad una interessante ricerca di produzione di due Higgs nel canale VBF.

1.3.1 Canale gluon-gluon fusion

ATLAS ha raccolto, tra il 2015 ed il 2016, un campione di dati corrispondente ad una luminosità integrata di 36.1 fb^{-1} ad una energia di $\sqrt{s} = 13 \text{ TeV}$. Le ricerche di coppie di Higgs sono state effettuate nell'ipotesi che i branching ratios e le proprietà cinematiche dell'Higgs siano quelle previste dal Modello Standard. Inoltre, sono stati individuati diversi sottocampioni corrispondenti ai possibili canali di decadimento:

- $HH \rightarrow b\bar{b}b\bar{b}$;
- $HH \rightarrow b\bar{b}W^+W^-$;
- $HH \rightarrow b\bar{b}\tau^+\tau^-$;
- $HH \rightarrow W^+W^-W^+W^-$;
- $HH \rightarrow b\bar{b}\gamma\gamma$;
- $HH \rightarrow W^+W^-\gamma\gamma$.

In tabella 1.4 sono riassunte le caratteristiche principali delle analisi relative ad ogni campione: branching ratio, discriminante utilizzato per l'individuazione di un eventuale segnale ($n.e.$ indica il numero di eventi), modello di produzione ricercato (NR indica una produzione risonante, S un bosone di Higgs scalare e G un gravitone) e massa della risonanza ipotizzata. Le quantità e i discriminanti indicati in parentesi quadre sono relativi alle ricerche di produzione risonante di coppie di Higgs.

Nei paragrafi successivi sono forniti dei dettagli sulle analisi effettuate in ciascun campione, prestando particolare attenzione ai tre canali che risultano essere più sensibili: $HH \rightarrow b\bar{b}b\bar{b}$, $HH \rightarrow b\bar{b}\tau^+\tau^-$, $HH \rightarrow b\bar{b}\gamma\gamma$. La ricostruzione dei diversi oggetti fisici (elettroni, muoni, leptoni τ , altre particelle cariche, fotoni e jet) è affidata a procedure complesse ma standardizzate che si basano sulla capacità di determinazione di momento, energia e natura dei diversi tipi di radiazione che caratterizzano i vari sottorivelatori di ATLAS, descritti in Appendice B.

Canale	BR	Discriminante	Modello	$m_{S/G}$ [TeV]
$b\bar{b}b\bar{b}$	0.34	m_{HH} [m_{HH}]	NR [S/G]	[0.26 – 3.00]
$b\bar{b}W^+W^-$	0.25	n.e. [m_{HH}]	NR [S/G]	[0.50 – 3.00]
$b\bar{b}\tau^+\tau^-$	0.073	BDT [BDT]	NR [S/G]	[0.26 – 1.00]
$W^+W^-W^+W^-$	0.046	n.e. [n.e.]	NR [S]	[0.26 – 0.50]
$b\bar{b}\gamma\gamma$	2.6×10^{-3}	$m_{\gamma\gamma}$ [m_{HH}]	NR [S]	[0.26 – 1.00]
$W^+W^-\gamma\gamma$	1.0×10^{-3}	$m_{\gamma\gamma}$ [$m_{\gamma\gamma}$]	NR [S]	[0.26 – 0.50]

Tabella 1.4: Caratteristiche principali delle analisi svolte su vari campioni, ognuno relativo ad un diverso canale di decadimento di coppie di bosoni di Higgs.

1.3.1.1 $HH \rightarrow b\bar{b}b\bar{b}$

Tale canale di decadimento rappresenta quello con il maggiore branching ratio, come indicato nella tabella 1.4. I dati raccolti sono stati analizzati utilizzando due diverse tecniche.

La prima analisi, detta *resolved analysis*, è adoperata nel caso in cui i due H abbiano un boost di Lorentz sufficientemente piccolo affinché i 4 b -jet prodotti possano essere ricostruiti. Con tale tecnica si ricerca una produzione non risonante oppure una risonante con massa della risonanza compresa tra 260 GeV e 1400 GeV. Il discriminante utilizzato per individuare la presenza di un segnale è la massa invariante dei 4 jet (e dunque dal sistema dei due Higgs, m_{HH}): la presenza di una risonanza è evidenziata da un eccesso localizzato, quella di una produzione non risonante da un eccesso nella coda dello spettro. Il fondo in questa analisi è costituito per il 95% da eventi multi-jet dovuti ad interazioni nucleari forti, per il 5% da eventi $t\bar{t}$, dal momento che t ($\bar{t}$) decade nel 100% dei casi in W^+b ($W^-\bar{b}$); il fondo proveniente da Z +jet o da processi che coinvolgono la produzione di un singolo Higgs è trascurabile.

La seconda tecnica utilizzata è detta *boosted analysis* ed è adoperata quando il boost di Lorentz dei due Higgs è così alto che i due b -jet generati da ogni decadimento non possono essere ricostruiti separatamente; pertanto si richiede l'individuazione di un singolo jet associato ad ogni Higgs all'interno del quale si ricerca la presenza di quark b . Gli eventi selezionati sono suddivisi in base al numero di b -tag osservati (la descrizione della procedura di b -tagging sarà oggetto del capitolo 2): il primo sottocampione include gli eventi per i quali ognuno dei due jet contiene un b -tag, il secondo quelli per i quali solo uno dei due jet contiene due b -tag ed il terzo quello per i quali entrambi i jet contengono 2 b -tag. La boosted analysis è utilizzata per la ricerca di risonanze di alta massa, compresa tra 800 GeV e 3000 GeV (nella regione di sovrapposizione con le masse delle risonanze cercate con la resolved analysis è stata utilizzata una combinazione delle due tecniche). Analogamente al caso precedente i principali processi di fondo sono la produzione di eventi multi-jet (responsabili dell' (80÷90)% del totale a seconda del sottocampione considerato) e gli eventi $t\bar{t}$ (responsabili del restante (10÷20)% del fondo); gli eventi Z +jet e quelli associati alla produzione del singolo Higgs sono trascurabili. Anche in questo caso il discriminante utilizzato per l'individuazione di un eventuale segnale è la massa invariante del sistema dei jet [12].

1.3.1.2 $HH \rightarrow b\bar{b}\tau^+\tau^-$

Al campione di eventi $HH \rightarrow b\bar{b}\tau^+\tau^-$ contribuiscono quelli per i quali almeno uno dei due τ decade adronicamente, mentre l'altro può decadere sia leptonicamente che adronicamente. In entrambi i casi la segnatura dell'evento è rappresentata dalla ricostruzione di due τ di carica opposta, di energia trasversa mancante dovuta ai neutrini provenienti dal decadimento dei τ e di due b -jet.

Gli eventi $\tau_{lep}\tau_{had}$ (per i quali uno dei due τ decade leptonicamente ed uno adronicamente) sono utilizzati per la ricerca di una produzione non risonante di coppie di Higgs e per quella di una risonanza di massa compresa tra 260 GeV e 800 GeV; il fondo principale in questo canale è rappresentato dalla produzione di $t\bar{t}$. Nel canale $\tau_{had}\tau_{had}$ (in cui entrambi i τ decadono adronicamente) le principali sorgenti di fondo sono rappresentate da eventi $t\bar{t}$, $Z \rightarrow \tau\tau$ ed eventi multi-jet. Questi eventi sono utilizzati per la ricerca di una produzione non risonante di coppie di Higgs e per quella di risonanza di massa compresa tra 260 GeV e 1000 GeV.

Una sorgente di fondo comune ad entrambi i canali è la presenza di finti τ_{had} , ovvero di jet originati da quark o gluoni ma erroneamente attribuiti ad un τ_{had} . Per distinguere l'eventuale presenza di un segnale dal fondo sono utilizzati algoritmi di *boosted decision tree*⁽⁵⁾ [13].

1.3.1.3 $HH \rightarrow b\bar{b}\gamma\gamma$

A questa categoria appartengono gli eventi con una segnatura costituita da due fotoni isolati e due jet con massa invariante compatibile con la massa dell'Higgs. Tali eventi sono suddivisi in sottocategorie distinte a seconda del fatto che solo uno oppure entrambi i jet siano taggati come b ; nel primo caso i jet devono soddisfare criteri sul momento trasverso e sulla massa invariante più stringenti rispetto al secondo.

La selezione più lasca comporta l'individuazione di un jet con $p_T > 40$ GeV e di uno con $p_T > 25$ GeV; è utilizzata per la ricerca di una risonanza scalare di massa compresa tra 260 GeV e 500 GeV. La selezione più stringente richiede che i due jet abbiano valori di p_T superiori a 100 GeV e 30 GeV rispettivamente; è sfruttata per la ricerca di una produzione non risonante e per quella di una risonanza scalare con massa compresa tra 500 GeV e 1000 GeV (tali eventi sono associati alla produzione di coppie di Higgs con momento trasverso maggiore rispetto a quelle prodotte da risonanze di bassa massa). I fotoni sono ricostruiti a partire da depositi energetici in celle contigue del calorimetro elettromagnetico (indicate nel seguito come *cluster*), ed è richiesto che non abbiano una corrispondenza con una traccia rivelata nell'Inner Detector⁽⁶⁾, caratteristica di un elettrone; la loro energia trasversa deve essere superiore a 35 GeV e 25 GeV rispettivamente.

Gli eventi $HH \rightarrow b\bar{b}\gamma\gamma$ sono affetti dal fondo proveniente da eventi con un singolo Higgs (in particolare nei canali $t\bar{t}H$ e ZH) e da quello dovuto a processi con stati finali caratterizzati da $\gamma\gamma$, γj e jj , nei quali i jet possono essere erroneamente ricostruiti come fotoni.

Il discriminante utilizzato per la ricerca di una produzione non risonante è la massa invariante dei due fotoni; nel caso della produzione risonante si utilizza la massa invariante del sistema costituito dai due fotoni e dai due jet [14].

⁽⁵⁾Un *boosted decision tree* è un algoritmo che combina i risultati di diversi *decision tree*, ciascuno dei quali suddivide i dati in ingresso in diverse categorie sulla base delle loro caratteristiche.

⁽⁶⁾Si veda l'Appendice B per la definizione del calorimetro e dell'Inner Detector.

1.3.1.4 Altri canali di decadimento

- $HH \rightarrow b\bar{b}W^+W^-$: le ricerche di ATLAS si sono concentrate in misura maggiore sul canale in cui una delle due W decade adronicamente, producendo dunque due quark, ed una leptonicamente, generando un leptone ed il neutrino ad esso associato; tale scelta deriva dal fatto la sua segnatura è più facilmente distinguibile dal fondo di QCD rispetto al caso in cui entrambi i bosoni decadono adronicamente. Inoltre, questo canale ha una statistica maggiore rispetto al caso in cui le W decadono leptonicamente (i primi due processi hanno un branching ratio del 45%, l'ultimo del 10%), sebbene questo sia il canale più pulito⁽⁷⁾. Pertanto, la segnatura ricercata è rappresentata dalla ricostruzione di un leptone, di elevata energia trasversa mancante e di almeno 4 jet.
Analogamente al caso $HH \rightarrow b\bar{b}b\bar{b}$, anche per questo canale di decadimento si utilizza la resolved analysis quando i due b -jet possono essere distinti tra loro e la boosted analysis quando il loro elevato boost di Lorentz non consente tale separazione. La resolved analysis è utilizzata per la ricerca di produzione non risonante di coppie di Higgs, di una risonanza scalare di massa compresa tra 500 GeV e 1400 GeV e di una di spin 2 con massa compresa tra 500 GeV e 800 GeV; la boosted analysis è adoperata per la ricerca di una risonanza scalare di massa compresa tra 1400 GeV e 3000 GeV e di una di spin 2 con massa compresa tra 800 e 3000 GeV [15].
- $HH \rightarrow W^+W^-W^+W^-$: è utilizzato per investigare la produzione non risonante di coppie di Higgs e quella risonante mediata dal bosone scalare neutro di massa maggiore previsto dalla teoria 2HDM; tale particella potrebbe decadere direttamente in due Higgs oppure in una coppia di nuovi bosoni scalari, ognuno dei quali decadrebbe successivamente in coppie W^+W^- con la stessa dipendenza del branching ratio dalla massa di quella del bosone di Higgs del Modello Standard. La massa della risonanza cercata è compresa tra 260 GeV e 500 GeV. Gli stati finali analizzati in questo campione sono quelli in cui 2, 3 o 4 W decadono leptonicamente [16].
- $HH \rightarrow W^+W^-\gamma\gamma$: sono inclusi gli eventi nei quali una W decade leptonicamente ed una adronicamente. Di conseguenza, la segnatura ricercata è costituita da due fotoni, due jet, un leptone carico ed un neutrino. Gli eventi $HH \rightarrow W^+W^-\gamma\gamma$ sono analizzati per la ricerca di una produzione non risonante di coppie di Higgs e per quella di una risonanza scalare di massa compresa tra 260 GeV e 500 GeV [17].

1.3.2 Canale vector boson fusion

Nel 2020 ATLAS ha pubblicato uno studio relativo alla produzione di coppie di Higgs attraverso il secondo canale di produzione più probabile: quello della vector boson fusion [18]. I dati analizzati sono stati raccolti tra il 2016 ed il 2018 e corrispondono ad una luminosità integrata di 126 fb^{-1} .

Come introdotto nella sezione 1.2.2, il canale VBF ha una dipendenza da k_λ inferiore rispetto a quella del canale gluon-gluon fusion. Per questo motivo il parametro della EFT rispetto al quale si studiano i dati raccolti è k_{2V} , che caratterizza l'interazione tra due bosoni di gauge V e due Higgs.

⁽⁷⁾Maggiori informazioni riguardo ai possibili decadimenti della coppia $t\bar{t}$ sono fornite nella sezione 2.4.

La segnatura della VBS è costituita da due jet di alto p_T , di massa invariante maggiore di 1000 GeV e con $|\Delta\eta| > 5$. Lo stato finale ricercato è quello in cui entrambi gli Higgs decadono nel modo più probabile: $HH \rightarrow b\bar{b}b\bar{b}$. Le due principali sorgenti di fondo sono rappresentate da eventi multi-jet (che rappresentano il 95% del totale) e dalla produzione di $t\bar{t}$ (5% del totale). Il discriminante utilizzato per individuare la presenza di un eventuale segnale è la massa invariante dei 4 b -jet, m_{4b} . Il confronto tra i dati raccolti, il fondo atteso e le predizioni relative alla produzione non risonante e a quella di una risonanza stretta di massa 800 GeV è rappresentato a sinistra in figura 1.9; nel pannello inferiore è rappresentato il rapporto tra i dati ed il fondo. Non si osserva alcun eccesso significativo⁽⁸⁾ rispetto al fondo: la più grande deviazione, visibile

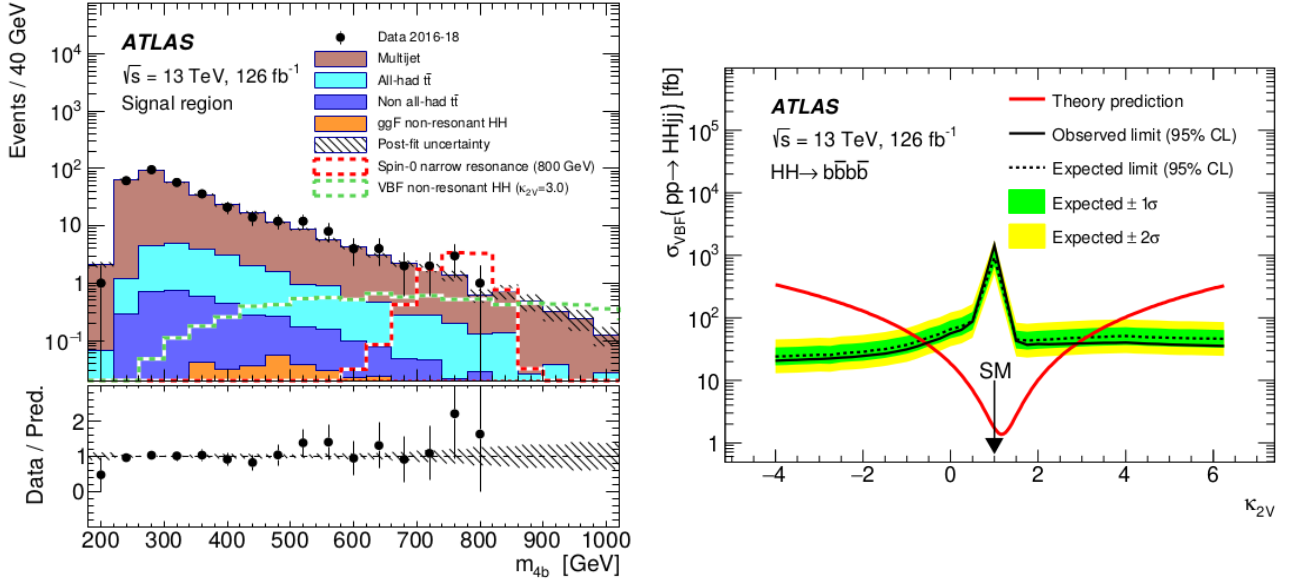


Figura 1.9: Sinistra: confronto tra i dati osservati, il fondo atteso e due diverse ipotesi di segnale in funzione della massa invariante dei 4 b -jet. Nel pannello inferiore è rappresentato il rapporto tra i dati ed il fondo. Destra: limiti superiori al livello di confidenza del 95% sulla sezione d'urto di produzione non risonante di coppie di Higgs in funzione del valore di k_{2V} . Da [18].

a 550 GeV, è associata ad una significanza di 1.5σ .

I dati raccolti hanno permesso di porre dei limiti superiori al livello di confidenza del 95% sulla sezione d'urto di produzione non risonante di coppie di Higgs in funzione del valore di k_{2V} ; il risultato ottenuto è rappresentato in figura 1.9 (destra). Gli intervalli consentiti per k_{2V} al livello di confidenza del 95% sono:

$$k_{2V} < -0.56 \cup k_{2V} > 2.89 \quad \text{osservato}$$

$$k_{2V} < -0.67 \cup k_{2V} > 3.10 \quad \text{atteso.}$$

Inoltre, sono stati individuati dei limiti superiori alla sezione d'urto di produzione di una generica particella di massa M che decada in una coppia di Higgs con branching ratio $\sim 100\%$: per $M \sim 1$ TeV tale limite è ~ 3 fb nel caso in cui la risonanza abbia una larghezza trascurabile

⁽⁸⁾ La significanza necessaria per dichiarare un'evidenza di segnale è pari a 3σ , mentre una scoperta richiede una significanza superiore a 5σ .

rispetto alla risoluzione sperimentale e ~ 8 fb in caso contrario, per $M \sim 200$ GeV il limite è ~ 1 pb indipendentemente dalla larghezza della risonanza.

1.4 Interpretazione dei risultati: produzione non risonante

Come già accennato, la ricerca di produzione non risonante di coppie di Higgs è effettuata assumendo che i branching ratios e le proprietà cinematiche dell'Higgs siano quelli previsti dal Modello Standard, mentre la sezione d'urto del processo di produzione di coppie di Higgs tramite gluon-gluon fusion (σ_{ggF}) può variare rispetto a quella del Modello Standard. In figura 1.10 (sinistra) sono rappresentati i limiti superiori al livello di confidenza del 95% sul rapporto $\sigma_{ggF}/\sigma_{ggF}^{SM}$. Si osserva che i canali di decadimento che determinano i limiti più stringenti sono $HH \rightarrow b\bar{b}\tau^+\tau^-$, $HH \rightarrow b\bar{b}b\bar{b}$ e $HH \rightarrow b\bar{b}\gamma\gamma$, come anticipato nella sezione 1.3.1. Il limite

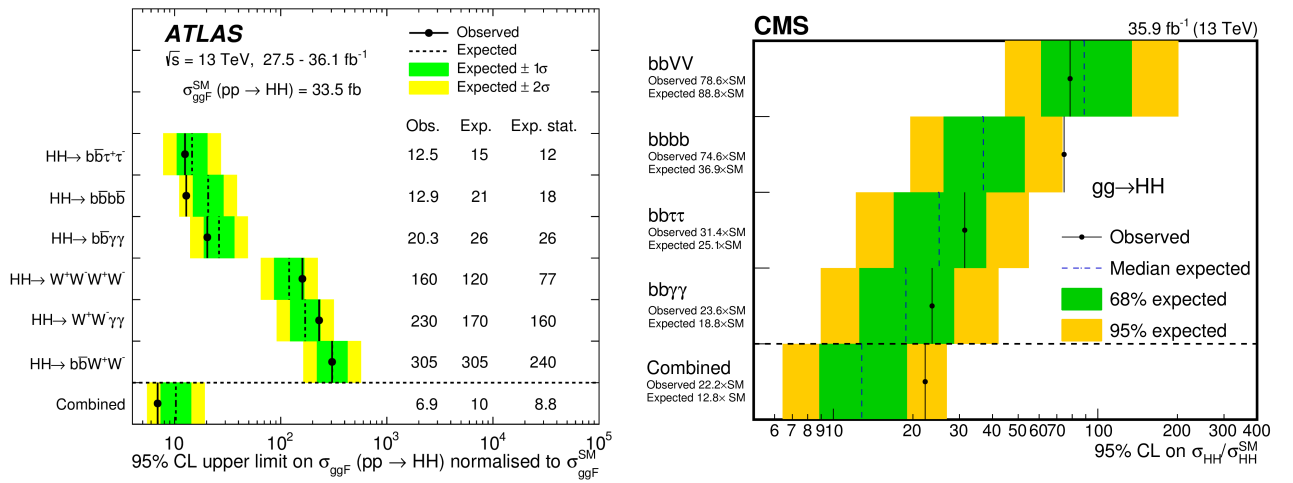


Figura 1.10: Limiti superiori al livello di confidenza del 95% sul rapporto tra la sezione d'urto di produzione di coppie di Higgs tramite gluon-gluon nell'ipotesi di fisica BSM e la sezione d'urto prevista dal Modello Standard per i diversi canali decadimento analizzati da ATLAS (sinistra) e CMS (destra). Da [19] e [20] rispettivamente.

atteso è pari a $10 \times \sigma_{ggF}^{SM}$, ed è compatibile entro 2σ con quello osservato, $6.9 \times \sigma_{ggF}^{SM}$. In figura sono riportati anche i valori attesi nel caso in cui si considerino i soli errori statistici, isolando l'impatto degli errori sistematici, dovuti alla modellizzazione del fondo, alla simulazione degli eventi ed alla identificazione e ricostruzione dei τ .

In figura 1.10 (destra) sono riportati i risultati ottenuti da un'analogia analisi effettuata da CMS nella quale ai canali $HH \rightarrow b\bar{b}VV$, $HH \rightarrow b\bar{b}b\bar{b}$, $HH \rightarrow b\bar{b}\tau^+\tau^-$, $HH \rightarrow b\bar{b}\gamma\gamma$ si aggiunge la combinazione di canali in cui un Higgs decade in $b\bar{b}$ e l'altro in una coppia di bosoni di gauge. I dati analizzati sono stati raccolti durante il 2016 e corrispondono ad una luminosità integrata di 35.9 fb⁻¹. Si osserva che il limite atteso è pari a $12.8 \times \sigma_{ggF}^{SM}$, quello osservato è $22.2 \times \sigma_{ggF}^{SM}$; anche in questo caso i due valori sono compatibili tra loro entro 2σ .

I dati relativi ai canali $HH \rightarrow b\bar{b}\tau^+\tau^-$, $HH \rightarrow b\bar{b}b\bar{b}$, $HH \rightarrow b\bar{b}\gamma\gamma$, caratterizzati da una maggiore sensibilità, sono stati utilizzati da ATLAS per studiare la compatibilità degli eventi selezionati con diverse ipotesi sul valore di k_λ . Questo parametro, intervenendo nel calcolo dell'ampiezza del diagramma in alto a sinistra in figura 1.4, influisce sulle proprietà cinematiche

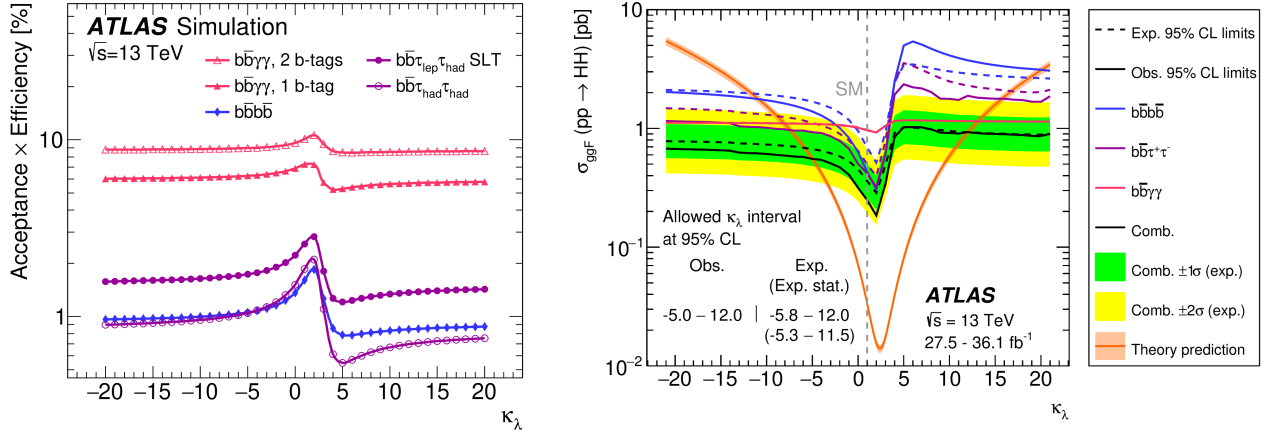


Figura 1.11: Sinistra: prodotto tra l'accettanza e l'efficienza in funzione di κ_λ . Destra: limite superiore al livello di confidenza del 95% sulla sezione d'urto di produzione non risonante di coppie di Higgs a partire da un processo di gluon-gluon fusion in funzione di κ_λ . La linea verticale indica l'ipotesi in cui κ_λ assume il valore predetto dal Modello Standard. Da [19].

dell'evento considerato e, pertanto, determina una modulazione nel prodotto di accettanza ed efficienza della selezione illustrato a sinistra in figura 1.11. Il massimo localizzato intorno a $\kappa_\lambda = 2.4$ corrisponde alla massima interferenza distruttiva tra i primi due diagrammi in figura 1.6, associata alla regione in cui gli Higgs ed i loro prodotti di decadimento hanno energie mediamente più elevate e dunque molti eventi superano i tagli di selezione; il minimo delle curve in figura è associato invece alla regione in cui tali particelle hanno energie mediamente più basse.

I risultati appena descritti possono essere sfruttati per porre dei limiti superiori al valore della sezione d'urto in funzione di κ_λ ; i risultati ottenuti sono rappresentati a destra in figura 1.11, nella quale è raffigurato un confronto con le predizioni teoriche. È possibile individuare una corrispondenza tra i valori di κ_λ per i quali si osservano un minimo ed un massimo nella figura precedente e quelli in corrispondenza dei quali si verificano un massimo ed un minimo nella figura in esame. Dall'analisi effettuata risulta che gli intervalli ammessi per κ_λ al livello di confidenza del 95% sono:

$$\begin{aligned} -5.0 < \kappa_\lambda < 12.0 & \quad \text{osservato} \\ -5.8 < \kappa_\lambda < 12.0 & \quad \text{atteso.} \end{aligned} \tag{1.2}$$

Il valore di κ_λ influenza, oltre la sezione d'urto di produzione di coppie di Higgs, anche i processi di produzione non risonante e di decadimento di un singolo Higgs attraverso le correzioni elettrodeboli al NLO (maggiori informazioni possono essere trovate nell'Appendice A); tale dipendenza non è stata considerata nel derivare i risultati appena descritti, ma è stato calcolato che la sua influenza percentuale massima è del 7%.

La dipendenza della sezione d'urto da κ_λ è stata analizzata anche da CMS utilizzando gli stessi canali di decadimento analizzati in figura 1.11 (destra); i suoi risultati sono riportati in figura 1.12. Da tale analisi risulta che gli intervalli di variabilità consentiti al livello di confidenza del 95% per κ_λ sono:

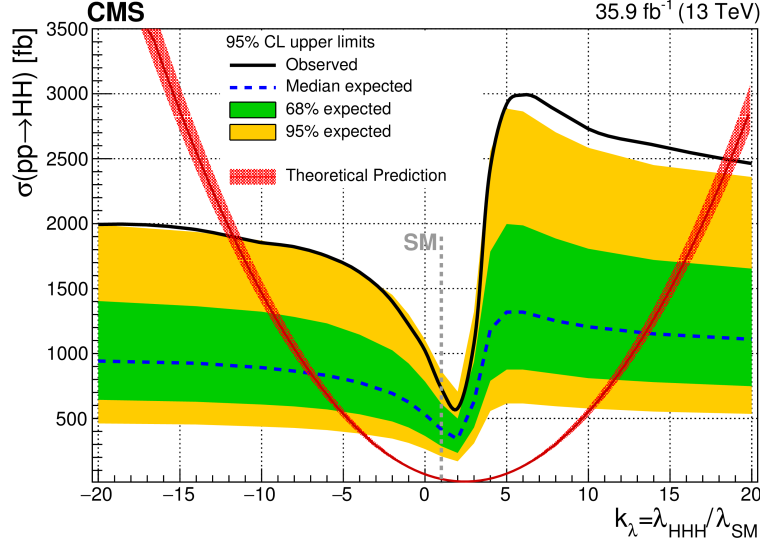


Figura 1.12: Limite superiore alla sezione d'urto di produzione non risonante di coppie di Higgs in funzione di k_λ . La linea verticale indica l'ipotesi k_λ assume il valore predetto dal Modello Standard. Da [20].

$$\begin{aligned} -11.8 < k_\lambda < 18.8 & \quad \text{osservato} \\ -7.1 < k_\lambda < 13.6 & \quad \text{atteso.} \end{aligned}$$

1.5 Interpretazione dei risultati: produzione risonante

Come illustrato in tabella 1.4 i sei possibili canali di decadimento delle coppie di bosoni di Higgs prodotte tramite gluon-gluon fusion sono utilizzati per ricercare una risonanza scalare di massa compresa tra 260 GeV e 3000 GeV, mentre l'esistenza di un gravitone di spin 2 è investigata sfruttando i canali di decadimento più significativi: $HH \rightarrow b\bar{b}b\bar{b}$, $HH \rightarrow b\bar{b}W^+W^-$ e $HH \rightarrow b\bar{b}\tau^+\tau^-$.

In figura 1.13 sono rappresentati i limiti superiori al livello di confidenza 95% sulla sezione d'urto di produzione risonante in funzione della massa della risonanza scalare ipotizzata; le linee verticali separano gli intervalli in cui diversi stati finali sono stati combinati tra loro. La maggiore deviazione dal fondo osservata è di 1σ , e pertanto non è significativa.

In figura 1.14 sono illustrati i limiti superiori alla produzione risonante di coppie di Higgs a partire da un gravitone in funzione della sua massa; i colori utilizzati fanno riferimento alla legenda in figura 1.13. I due grafici in figura 1.14 sono relativi a due diverse ipotesi dei parametri k e $\bar{M}_{\text{Pl}}$ che caratterizzano la teoria: k è legato alla curvatura dello spazio-tempo multidimensionale, $\bar{M}_{\text{Pl}} = 2.4 \times 10^{18}$ GeV è la massa di Planck. Il risultato a sinistra si riferisce all'ipotesi $k/\bar{M}_{\text{Pl}} = 1$, quello a destra all'ipotesi $k/\bar{M}_{\text{Pl}} = 2$. In nessun caso si osserva una deviazione significativa dal fondo: nel primo la discrepanza maggiore ha una significatività di 1.5σ , nella seconda di 0.7σ . Questi risultati hanno permesso di escludere al livello di confidenza del 95% la presenza di un gravitone con parametri tali che $k/\bar{M}_{\text{Pl}} = 1$ nell'intervallo di masse compreso tra 310 GeV e 1380 GeV e quella di un gravitone con parametri tali che $k/\bar{M}_{\text{Pl}} = 2$ e massa nell'intervallo compreso tra 260 GeV e 1760 GeV.

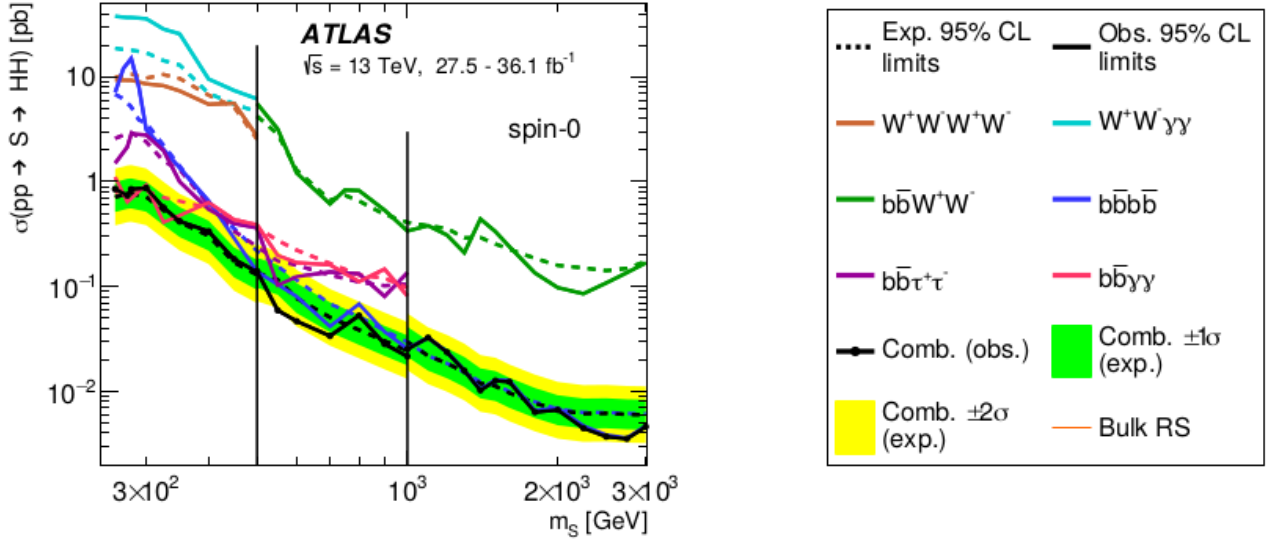


Figura 1.13: Limiti superiori al livello di confidenza del 95% sulla sezione d'urto di produzione risonante in funzione della massa della risonanza scalare ipotizzata. Le linee verticali indicano gli intervalli in cui sono stati combinati gli stessi stati finali. Da [19].

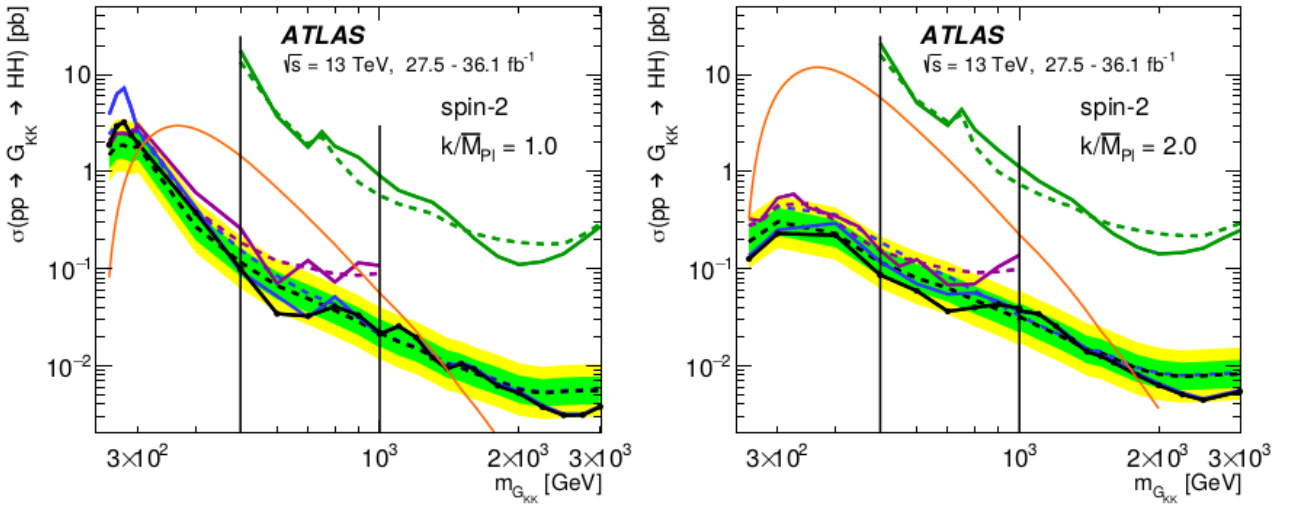


Figura 1.14: Limiti superiori relativi alla produzione risonante di coppie di Higgs a partire da un gravitone in funzione della massa della risonanza nel caso $k/\overline{M}_{Pl} = 1$ (sinistra) e $k/\overline{M}_{Pl} = 2$ (destra). I colori utilizzati fanno riferimento alla legenda in figura 1.13; le linee verticali indicano gli intervalli in cui sono stati combinati diversi stati finali. Da [19].

Analoghi risultati sono stati ottenuti dalla collaborazione CMS [20].

Capitolo 2

Il b -tagging ad ATLAS

Nella ricerca di produzione di coppie di bosoni di Higgs i canali più sensibili sono quelli che presentano coppie $b\bar{b}$ nello stato finale. Ciò implica la necessità di identificare tali quark con la maggiore precisione possibile, in modo da poterli distinguere dal fondo. L'individuazione dei quark b è una tematica ricorrente anche in numerose analisi precedenti alla più recente ricerca di coppie HH . Gli algoritmi utilizzati da ATLAS per raggiungere tale obiettivo hanno subito varie modifiche volte a migliorarne le prestazioni e ad adeguarli ai cambiamenti apportati al rivelatore di tracciamento nel corso degli anni.

In questo capitolo sono introdotte le signature associate ai quark b e la logica seguita per la loro identificazione. Successivamente sono discusse le principali caratteristiche degli algoritmi utilizzati e l'impatto della geometria del rivelatore su di esse, attraverso un confronto tra le prestazioni ottenute durante il run 1 ed il run 2 e tra queste ultime e quelle previste durante il run di HL—LHC. Infine, sono illustrate l'estrapolazione dei risultati relativi alla produzione di coppie HH tramite gluon-gluon fusion alla luminosità di HL—LHC e le più recenti evoluzioni degli algoritmi utilizzati da ATLAS per l'identificazione dei b -jet.

2.1 Segnatura dei b -jet

Il processo di formazione di jet a partire da un quark prende il nome di *frammentazione*. Esso ha inizio con l'irraggiamento di gluoni da parte del quark, i quali generano uno sciame partonico al cui interno le particelle interagiscono dapprima lievemente, dato il basso valore della costante di accoppiamento forte ad alti momenti trasferiti. All'aumentare della separazione tra le particelle, tuttavia, il valore di tale costante aumenta e l'interazione tra i partoni è tale che ci sia energia sufficiente per creare coppie di quark dal vuoto. Il ripetersi di questo processo determina la formazione di adroni (fenomeno noto come *adronizzazione*). Questi hanno la componente dominante del momento nella direzione del quark originario, perciò appaiono come getti (*jet*) collimati di adroni, in uno dei quali è legato il quark prodotto nell'interazione. L'adronizzazione è un processo stocastico e non perturbativo; ciò fa sì che la ricostruzione di un jet sia inevitabilmente affetta da incertezze derivanti dal modello utilizzato per la sua descrizione.

Un jet creato dalla frammentazione di un quark b prende il nome di b -jet; la sua identificazione trae vantaggio da alcune proprietà di tale quark. Esso ha un'elevata funzione di frammentazione, la quale parametrizza la quantità di energia condivisa tra il quark originario e l'adrone che lo contiene nello stato finale: un adrone contenente un quark b ha in media l'80% della sua energia. Inoltre, il tempo di decadimento di un adrone con numero quantico di *beauty* diverso

da 0 (indicato nel seguito come *adrone con beauty*) è dell'ordine del ps. Tali proprietà fanno sì che il vertice secondario, dovuto al decadimento di un adrone contenente b , sia distanziato da quello primario, originato dall'interazione tra i protoni dei due fasci di LHC. La distanza media x percorsa da una particella di massa m , energia E , momento $\vec{p}$ e vita media τ prima di decadere è rappresentata da

$$x = \beta\gamma c\tau = \frac{|\vec{p}|}{E} \frac{E}{m} c\tau = \frac{\sqrt{E^2 - m^2}}{m} c\tau$$

in cui β e γ rappresentano i fattori relativistici di Lorentz e c la velocità della luce. Si consideri, a titolo esemplificativo, un mesone B^+ ($u\bar{b}$) risultante dalla frammentazione di un quark b di energia 10 GeV (ipotesi ragionevole nel caso dei quark prodotti dall'interazione protone - protone ad LHC). L'energia del mesone sarà $E \sim 8$ GeV; la sua massa è $m(B^+) = (5279.34 \pm 0.12)$ MeV, la sua vita media $\tau = (1.638 \pm 0.004)$ ps [21]. Introducendo tali valori nella relazione precedente si ottiene $x \sim 0.56$ mm, che rappresenta una lunghezza di decadimento risolvibile con rivelatori ad alta precisione quali quelli basati su Silicio, che permettono di ottenere una risoluzione spaziale dell'ordine delle decine di μm ; per via delle loro eccellenti prestazioni essi sono utilizzati nel rivelatore di tracciamento interno di ATLAS.

Il quark b , con una massa di $4.18^{+0.03}_{-0.02}$ GeV [21], è il più pesante quark con possibilità di adronizzare (il quark t , l'unico con massa superiore al b , decade prima di farlo); ciò comporta che i prodotti del decadimento di adroni con beauty abbiano una massa invariante superiore a quella delle particelle originate dal decadimento di altri adroni. Infine, gli adroni contenenti b decadono producendo un numero medio di particelle cariche pari a 5 [?], superiore a quello che si osserva in altri decadimenti; un vertice di decadimento dal quale sono generate n tracce cariche è denominato n -prong.

Le proprietà elencate concorrono a determinare una segnatura peculiare del decadimento di adroni con beauty, rappresentata in figura 2.1. Le quantità e gli oggetti rappresentati in figura sono ricostruiti all'interno del sistema di coordinate destrorso utilizzato da ATLAS, in cui l'origine è collocata nel punto di interazione tra i fasci (spesso indicato come *beam spot* o *vertice primario*), l'asse z è diretto lungo la loro direzione di propagazione, l'asse x è orientato verso il centro della circonferenza di LHC e l'asse y verso l'alto. Per ogni traiettoria elicoidale descritta dalle particelle all'interno del campo magnetico di ATLAS è possibile individuare il punto la cui proiezione nel piano trasverso (x, y) è la più vicina all'origine, detto *perigee*. La distanza di tale punto dall'origine nel piano trasverso è indicata con d_0 ed è detta *parametro di impatto trasverso*, la sua coordinata z con z_0 . Gli angoli formati dalla tangente alla traiettoria nel perigee con gli assi z e x sono indicati con θ e ϕ rispettivamente. Il prodotto $z_0 \sin \theta$ rappresenta il *parametro di impatto longitudinale*. È possibile inoltre definire la pseudorapidità come $\eta = -\ln \tan(\theta/2)$ e la distanza angolare come $\Delta R = \sqrt{(\Delta\eta)^2 + (\Delta\phi)^2}$. In figura 2.1 si osserva come l'adrone decada dopo aver percorso una distanza L_{xy} generando un vertice secondario ben distanziato da quello primario. Ciò fa sì che le tracce originate dal suo decadimento abbiano elevati parametri di impatto trasversi e longitudinali.

Un'altra particolarità degli adroni con beauty è rappresentata dal fatto che possono dare origine ad una catena di decadimenti in cui si distinguono due vertici secondari. Ciò accade nel caso in cui decadano debolmente in adroni contenenti c , i quali hanno anch'essi una vita media dell'ordine del ps e parametri di impatto elevati; pertanto il loro decadimento origina un vertice distanziato da quello associato al decadimento dell'adrone contenente b dal quale provengono.

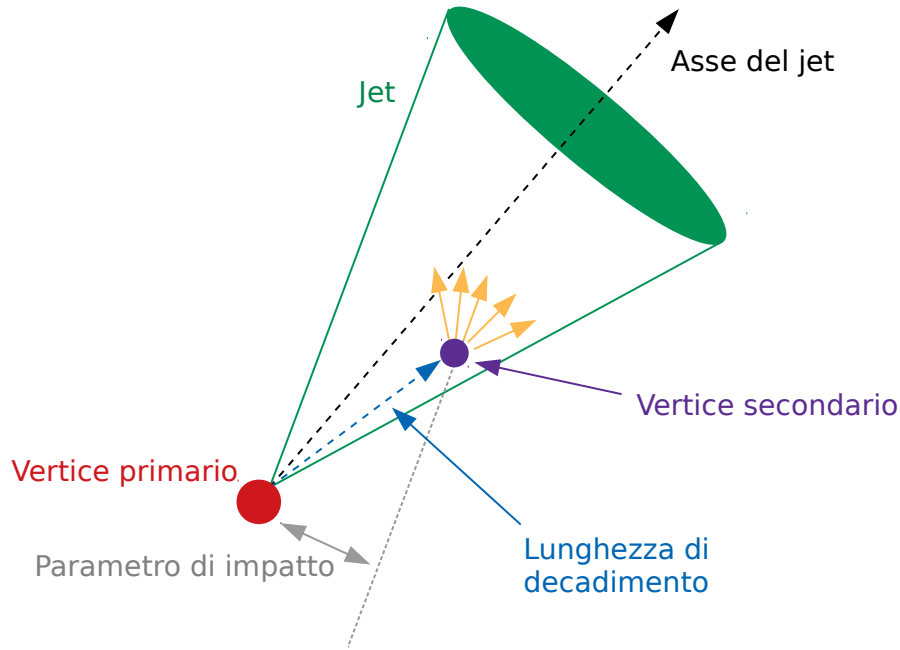


Figura 2.1: Tipica segnatura associata al decadimento di un adrone contenente un quark b generato dall'interazione protone - protone ad LHC.

2.2 Oggetti fisici utilizzati per il b -tagging

L'identificazione dei b -jet, processo noto come b -tagging, richiede la ricostruzione di tracce e jet.

Le tracce sono ricostruite raccordando i punti misurati dall'Inner Detector, detti *hit*. Dal fit della traiettoria determinata dal campo magnetico alla sequenza di hit si determina il momento della particella. Le tracce sono utilizzate per l'individuazione dei vertici. In corrispondenza di ogni incrocio dei fasci si verificano diverse interazioni, di numero variabile e crescente all'aumentare della luminosità. Queste interazioni simultanee, indicate come *pile-up*, determinano l'individuazione di diversi vertici all'interno di ogni evento; il vertice primario rispetto al quale si svolge la sua analisi è quello per il quale la somma dei momenti trasversi delle tracce che contribuiscono alla sua ricostruzione è massima.

Parallelamente alla ricostruzione delle tracce avviene quella dei jet, realizzata a partire dall'identificazione di cluster energetici all'interno dei calorimetri adronici. Successivamente essi sono sottoposti ad una selezione, effettuata sulla base del loro momento trasverso e della loro pseudorapidità, finalizzata alla reiezione di finti jet. Un'importante sorgente di fondo nella ricostruzione dei jet è rappresentata dal *pile-up* e dall'osservazione di depositi energetici all'interno del calorimetro dovuti ad interazioni verificatesi in collisioni precedenti o successivi rispetto a quello contenente l'evento considerato; tali fenomeni sono talvolta indicati come *out-of-time pile-up*. La soppressione del *pile-up* avviene attraverso algoritmi dedicati, tra i quali ad esempio il *Jet Vertex Tagger* [23].

Una volta ricostruite le tracce e i jet questi sono associati tra loro imponendo dei limiti

superiori alla loro distanza angolare; tale limite è più stringente nel caso di jet ad alto p_T dal momento che i prodotti del decadimento di adroni contenenti *b* con alto momento trasverso sono più collimati rispetto a quelli originati dal decadimento di adroni a basso p_T . Nel caso in cui diversi jet soddisfino i criteri di associazione ad una traccia viene selezionato quello più vicino ad essa.

2.3 Valutazione dell'efficienza degli algoritmi di *b*-tagging tramite simulazione Monte Carlo

Le tracce che soddisfano i criteri di associazione ai jet sono utilizzate come input degli algoritmi di *b*-tagging, illustrati nella sezione 2.5. Questi hanno come obiettivo la determinazione della probabilità che un jet sia di tipo *b*, *c* oppure *light-flavour* (se generato da un quark di massa minore) sulla base di discriminanti statistici che sfruttano le caratteristiche tipiche di un *b*-jet, ovvero la presenza di vertici secondari distanziati dal primario e parametri di impatto elevati. Tali algoritmi, pertanto, forniscono come output le distribuzioni dei discriminanti alle quali sono imposti dei tagli volti ad isolare i *b*-jet dal fondo con una efficienza stabilita a priori, indicata come *punto di lavoro*. Le loro prestazioni sono valutate eseguendo delle simulazioni Monte Carlo, nelle quali la natura dei jet è nota. Questi infatti sono etichettati come *b*-jet se sono associati ad adroni contenenti *b* che soddisfano requisiti di prossimità, ovvero interni ad un cono con asse coincidente con quello del jet ed apertura tale da contenere una grande frazione della sua energia. L'informazione sulla natura e la posizione degli adroni è fornita dalla *verità Monte Carlo*, ossia dall'elenco delle particelle che sono state prodotte nell'evento a partire dall'emulazione dell'interazione primaria e fino allo sviluppo dei jet nei processi di frammentazione e adronizzazione. Nel caso in cui non siano identificati adroni contenenti *b* che soddisfino tale criterio di prossimità si ricercano quelli contenenti *c*, eventualmente etichettando il jet come *c*-jet; se anche tale individuazione fallisce si ricercano leptoni $\tau^{(1)}$, etichettando in tal caso il jet come τ -jet. Nel caso in cui le ricerche precedenti diano esito negativo il jet ricostruito è etichettato come *light-flavour jet*.

Le prestazioni degli algoritmi sono generalmente quantificate sulla base di due fattori: l'efficienza di *b*-tagging e la reiezione di jet di diverso flavour. L'*efficienza* è definita come il rapporto tra i jet che sono stati *taggati* come *b*-jet dagli algoritmi, ovvero quelli che hanno soddisfatto il taglio imposto sui loro output, e i *b*-jet presenti nel campione analizzato, ovvero quelli etichettati come *b*-jet sulla base della conoscenza della verità Monte Carlo e dei quali i primi rappresentano un sottoinsieme. La *reiezione* di *c*-jet è invece definita come l'inverso della probabilità di taggare erroneamente un *c*-jet come *b*-jet; una definizione analoga vale per la reiezione di *light-flavour jet*.

2.4 Calibrazione degli algoritmi di *b*-tagging sui dati

Le simulazioni Monte Carlo sono sfruttate per valutare la bontà di un algoritmo sulla base della procedura descritta nella sezione 2.3; tuttavia i risultati possono essere affetti da incertezze

⁽¹⁾Il leptone τ ha una vita media di $(290.3 \pm 0.5) \times 10^{-15}$ s, vicina a quella degli adroni con beauty, e decade in adroni con un branching ratio del 64.8% [21], dando vita a jet con caratteristiche simili a quelle dei jet originati da quark *b* e *c*.

sistematiche derivanti dal fatto che le simulazioni non riproducono perfettamente le reali caratteristiche dei jet e dell'evento nel suo complesso. Per questo motivo si applica una procedura di calibrazione degli algoritmi sui dati. A tale scopo si seleziona un campione puro di b -jet; questo viene da processi del tipo $t\bar{t}$, molto abbondanti ad LHC⁽²⁾, in cui t e $\bar{t}$ decadono dopo un tempo di 4×10^{-25} s [8] rispettivamente in W^+b e $W^-\bar{b}$, con un branching ratio di $\sim 100\%$. I bosoni prodotti possono decadere in coppie di quark o di leptoni, rendendo possibili tre diversi stati finali a cui corrispondono le segnature rappresentate in figura 2.2:

$$\begin{aligned} t\bar{t} &\rightarrow W^+bW^-\bar{b} \rightarrow q\bar{q}'bq''\bar{q}'''\bar{b} & (\text{BR} = 45.7\%) \\ t\bar{t} &\rightarrow W^+bW^-\bar{b} \rightarrow q\bar{q}'bl^-\bar{\nu}_l\bar{b} + l^+\nu_lbq''\bar{q}'''\bar{b} & (\text{BR} = 43.8\%) \\ t\bar{t} &\rightarrow W^+bW^-\bar{b} \rightarrow l^+\nu_lb l'^-\bar{\nu}_{l'}\bar{b} & (\text{BR} = 10.5\%) \end{aligned}$$

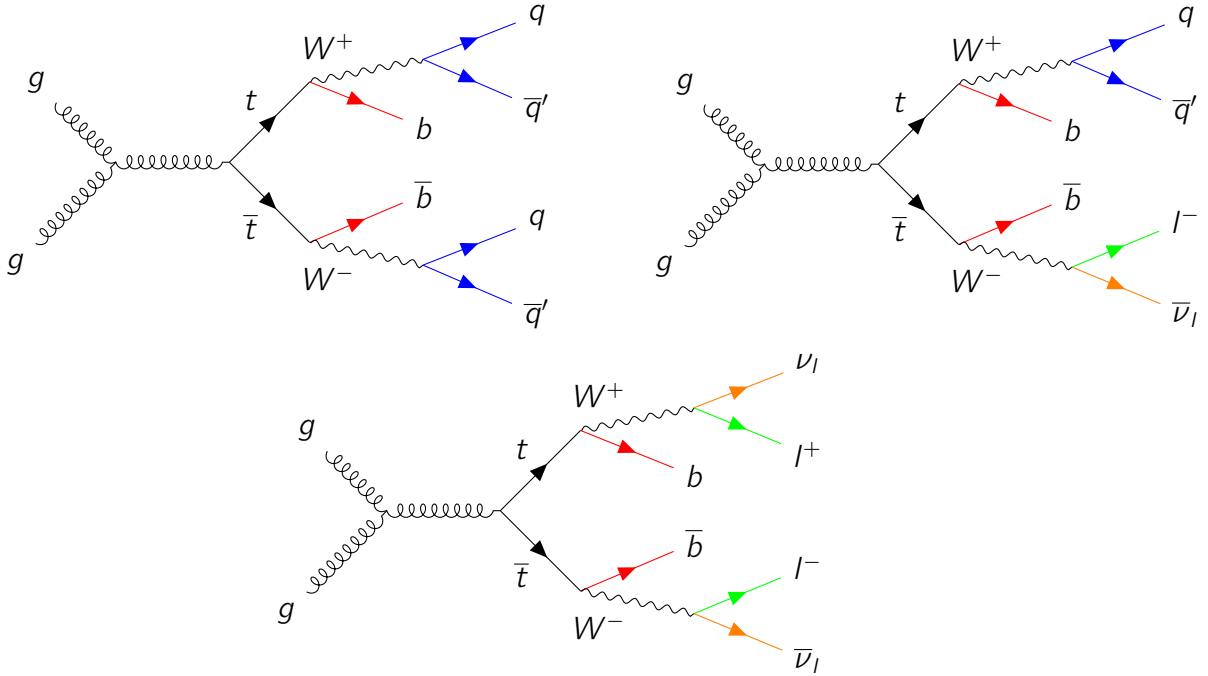


Figura 2.2: Possibili stati finali del decadimento della coppia $t\bar{t}$, prodotta ad LHC dall'interazione gluone - gluone (rappresentata in figura) o quark - antiquark (in misura minore dal momento che nel processo pp il contributo di antiquark proviene dal mare).

È evidente come nell'ultimo caso, indicato come *canale dileptonico*, gli unici jet nello stato finale siano di tipo b ; pertanto è possibile selezionare un campione puro di b -jet ricercando la segnatura associata all'ultimo diagramma in figura 2.2. L'elevata sezione d'urto di produzione di $t\bar{t}$ ad LHC permette di bilanciare il basso branching ratio del decadimento dileptonico delle W .

Una volta selezionato il campione di b -jet è possibile applicarvi gli algoritmi di b -tagging ed imporre ai loro output gli stessi tagli utilizzati nel caso della simulazione Monte Carlo. Il rapporto tra i jet taggati come b dagli algoritmi e quelli contenuti nel campione fornisce l'efficienza di

⁽²⁾Le collaborazioni ATLAS e CMS hanno calcolato una sezione d'urto di produzione di coppie $t\bar{t}$ pari a $\sigma_{t\bar{t}} = (803 \pm 2(\text{stat}) \pm 25(\text{sist}) \pm 20(\text{lumi}))$ pb utilizzando i dati raccolti all'energia $\sqrt{s} = 13$ TeV corrispondenti ad una luminosità integrata di 36 fb^{-1} [24]. L'errore indicato come *lumi* è associato alla misura della luminosità.

b -tagging. Dal rapporto tra l'efficienza ricavata sui dati (ϵ_{dati}) e quella ottenuta sulla base del Monte Carlo (ϵ_{MC}) è possibile calibrare gli algoritmi di b -tagging; il risultato di questa operazione è detto *fattore di scala*:

$$\text{fattore di scala (Scale Factor, SF)} = \frac{\epsilon_{\text{dati}}}{\epsilon_{\text{MC}}}.$$

Il fattore di scala è applicato alle simulazioni Monte Carlo per ripesare i jet taggati come b ed ottenere predizioni che riproducano fedelmente il comportamento degli algoritmi sui dati. Le discrepanze residue e gli errori sperimentali sui fattori di scala sono poi utilizzati per quantificare l'incertezza sistematica sull'efficienza o la reiezione degli algoritmi. Tale procedura è applicata per calibrare alcuni punti di lavoro prestabiliti, che per il run 2 di ATLAS corrispondono ad efficienze pari a 60%, 70%, 77% e 85%.

2.5 Algoritmi per l'identificazione di b -jet

Gli algoritmi utilizzati per l'identificazione di b -jet ad ATLAS possono essere distinti in due categorie: quelli di primo livello, che sfruttano le signature dei b -jet descritte nella sezione 2.1, e quelli di secondo livello, i quali combinano i risultati dei precedenti per ottenere prestazioni migliori. I grafici presentati nei prossimi paragrafi sono relativi alla simulazione di eventi $t\bar{t}$ ad un'energia nel centro di massa di 13 TeV, tra i quali sono selezionati solo quelli in cui almeno una delle W genera un leptone nello stato finale. Essi sono stati realizzati utilizzando le versioni degli algoritmi di b -tagging sfruttate nelle ricerche di coppie di bosoni di Higgs descritte nel capitolo 1.

2.5.1 Algoritmi di primo livello

Gli algoritmi di primo livello possono essere a loro volta suddivisi in diverse categorie a seconda delle strategie adottate: una prima tipologia sfrutta i parametri di impatto delle singole tracce associate al jet, mentre altri algoritmi combinano tali tracce per ricostruire esplicitamente i vertici secondari.

Algoritmi basati sul parametro di impatto: IP2D, IP3D

Come conseguenza della loro elevata vita media e del boost con il quale gli adroni contenenti quark b sono prodotti, le tracce generate dal loro decadimento sono caratterizzate da un elevato parametro di impatto, sia trasverso che longitudinale. Poiché il vertice di decadimento di tali adroni deve necessariamente giacere lungo la loro linea di volo si associa ai parametri di impatto un segno in modo da distinguere su base statistica le tracce che provengono da un vertice secondario da quelle originate dal vertice primario, indicate come tracce *prompt*. Tale segno è definito positivo se la traccia interseca l'asse del jet nello stesso emisfero in cui questo si sta propagando, negativo altrimenti. Sulla base di questa convenzione le tracce che provengono da un vertice secondario hanno un parametro di impatto positivo, mentre per quelle *prompt*, che dovrebbero avere un parametro di impatto nullo, la risoluzione intrinseca del rivelatore genera un parametro di impatto con segno che può essere positivo o negativo con uguale probabilità.

L'algoritmo *IP2D* [25] sfrutta la significanza sul parametro di impatto trasverso, definita come il rapporto tra d_0 e la sua risoluzione:

$$\frac{d_0}{\sigma(d_0)}.$$

L'algoritmo *IP3D* [25] utilizza anche la significanza sul parametro di impatto longitudinale, definita in maniera analoga come

$$\frac{z_0 \sin \theta}{\sigma(z_0 \sin \theta)}$$

e tiene conto della correlazione tra le due utilizzando template bidimensionali.

In figura 2.3 sono rappresentate le distribuzioni delle significanze dei parametri di impatto trasverso (sinistra) e longitudinale (destra) di tracce appartenenti a b -jet (in verde), c -jet (in blu) e light-flavour jet (in rosso). In tutte le curve è evidente l'effetto della risoluzione del

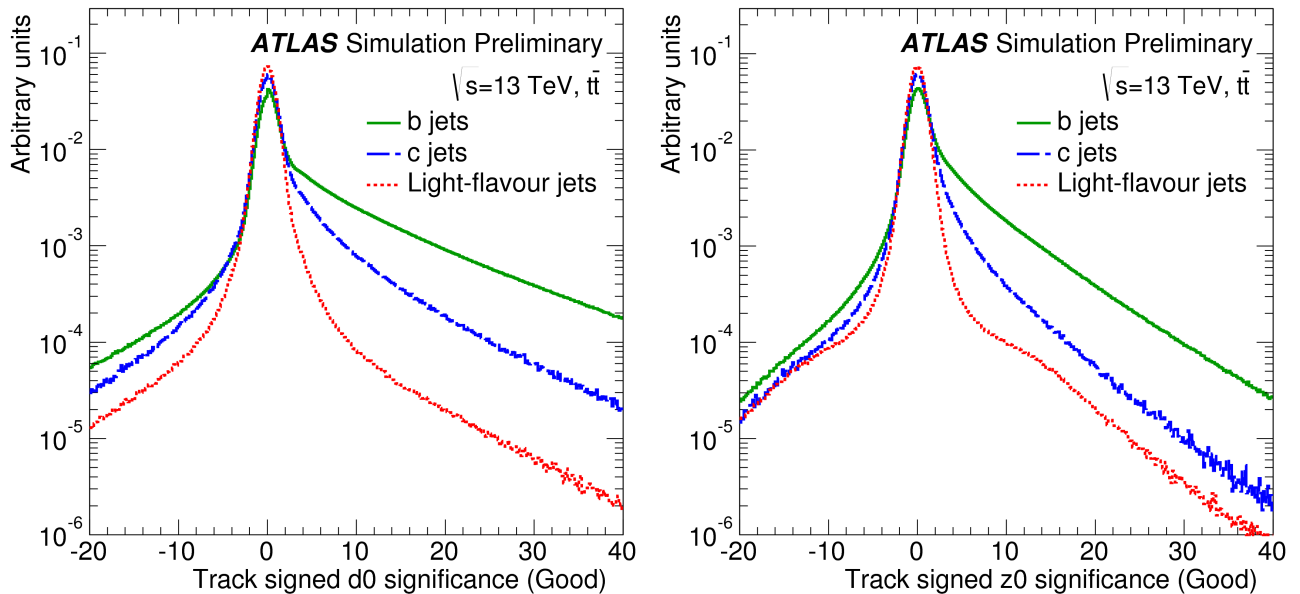


Figura 2.3: Distribuzione della significanza del parametro di impatto trasverso (sinistra) e longitudinale (destra) per b -jet (in verde), c -jet (in blu) e light-flavour jet (in rosso). Da [26].

rivelatore, che determina una componente simmetrica intorno allo 0. Nel caso dei light-flavour jet la coda leggermente più pronunciata ad alti valori positivi di significanza è attribuibile alle tracce provenienti dal decadimento di K_S e Λ , dalla conversione di fotoni e dall'interazione delle particelle con il materiale del rivelatore; nel caso di b -jet e c -jet la componente dominante ad alti parametri di impatto è dovuta al decadimento di adroni contenenti b o c . Nelle distribuzioni relative al parametro di impatto longitudinale si osserva come la possibilità che light-flavour jet siano erroneamente considerati come b -jet sia maggiore rispetto al caso in cui si consideri il solo parametro di impatto trasverso; ciò è una conseguenza del fatto che le tracce provenienti da pile-up erroneamente associate al jet hanno valori di z_0 abbastanza elevati. Le distribuzioni in figura sono ottenute per 14 categorie di tracce. La distinzione tra tali categorie è effettuata sulla base della topologia degli hit da esse determinati: si considera il numero di hit attesi in uno

o in entrambi gli strati più interni del rivelatore ma non osservati (informazione derivante dalla conoscenza del rivelatore), quelli *shared* (ovvero i cluster⁽³⁾ utilizzati per la ricostruzione di più tracce) e quelli *split* (cluster coincidenti dovuti alla sovrapposizione di più tracce). In figura 2.3 sono rappresentate le distribuzioni relative alla categoria indicata come *good*; essa comprende le tracce che soddisfano i più stringenti criteri di qualità e pertanto le curve raffigurate sono quelle con il maggior potere di reiezione di *c*-jet e light-flavour jet.

Le densità di probabilità (*Probabilities Densities Functions*, PDF) delle significanze così ottenute sono utilizzate per definire i discriminanti delle varie ipotesi di flavour dei jet. Le PDF forniscono la probabilità che la *i*-esima traccia sia associata ad un jet di un determinato flavour; le probabilità relative ai jet *b*, *c* e light-flavour sono indicate rispettivamente con $p_{b,i}$, $p_{c,i}$, $p_{u,i}$. Il discriminante utilizzato dagli algoritmi IP2D e IP3D per distinguere le ipotesi di *b*-jet e *c*-jet è definito come

$$\sum_{i=1}^N \log \left(\frac{p_{b,i}}{p_{c,i}} \right)$$

dove *N* rappresenta il numero di tracce associate ad un jet. Variabili analoghe sono utilizzate per distinguere le altre ipotesi di flavour. Le distribuzioni di tali discriminanti sono rappresentate in figura 2.4.

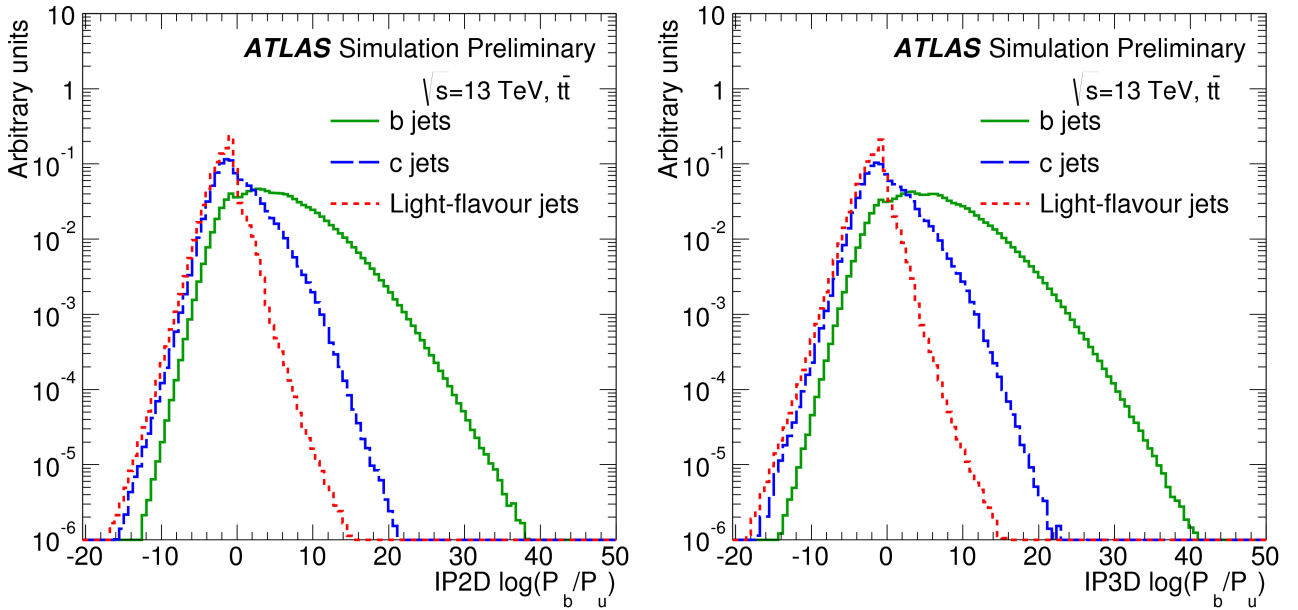


Figura 2.4: Distribuzione del discriminante risultante dall'algoritmo IP2D (sinistra) e IP3D (destra) per *b*-jet (in verde), *c*-jet (in blu) e light-flavour jet (in rosso). Da [26].

Algoritmo basato sulla ricostruzione di un vertice secondario: SV

L'algoritmo *Secondary Vertex 1* (SV1) [27] è utilizzato per la ricostruzione di un vertice secondario distanziato da quello primario. Le tracce associate ad un jet e sufficientemente distanti dal vertice primario nel punto di massimo avvicinamento sono utilizzate in coppia per ricostruire

⁽³⁾I cluster sono definiti all'interno del rivelatore di tracciamento come il raggruppamento di hit in pixel o strip contigui. Per la definizione di pixel e strip si veda la sezione 3.5.

possibili vertici secondari. Tra questi, quelli compatibili con l'interazione delle particelle con il materiale del rivelatore o con decadimenti abbondanti di π o K sono scartati. Le tracce associate ai vertici che hanno superato questa selezione sono utilizzate per ricostruire un singolo vertice secondario, applicandovi iterativamente l'algoritmo SV1; l'ipotesi di associazione al vertice secondario è valutata attraverso un test del χ^2 . Ad ogni step la traccia con il maggiore χ^2 è rimossa e la procedura termina nel momento in cui si individua un vertice secondario con un χ^2 accettabile ed una massa invariante minore di 6 GeV. Questo algoritmo associa ad un unico vertice secondario i prodotti di decadimento di adroni contenenti b o c .

La massa invariante del vertice ricostruito, la sua frazione di energia (definita come il rapporto tra l'energia delle tracce associate al vertice e quella delle tracce associate al jet) ed il numero di possibili vertici ottenuti dalla combinazione di coppie di tracce sono tre variabili di output utilizzate nel calcolo delle PDF relative all'ipotesi di b -jet, c -jet e light-flavour jet su cui si basano i discriminanti che definiscono le prestazioni dell'algoritmo. La massa invariante del vertice è calcolata a partire dalla somma dei quadrimpulsi delle tracce che contribuiscono alla determinazione del vertice. Poiché il tracciamento operato nell'Inner Detector fornisce informazioni unicamente sul momento delle particelle e non sulla loro natura è necessario avanzare delle ipotesi che consentano di determinare la loro energia e dunque il quadrimomento; l'ipotesi effettuata è che tutte le particelle provenienti dal vertice secondario abbiano massa uguale a quella del pione. In figura 2.5 sono rappresentate le distribuzioni di tali variabili per i b -jet, c -jet e light-flavour jet. Si osserva come la massa invariante ed il numero di possibili vertici siano

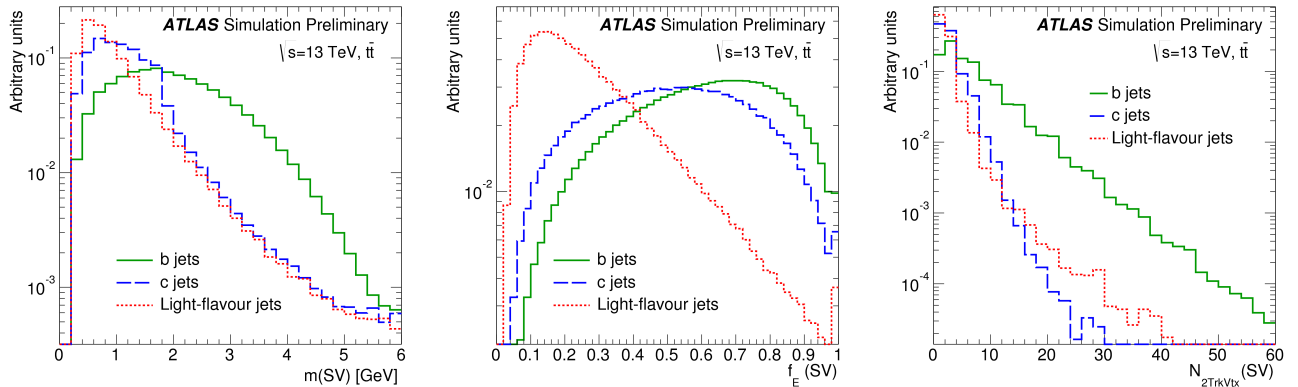


Figura 2.5: Distribuzione della massa invariante del vertice secondario ricostruito dall'algoritmo SV1 (sinistra), della sua frazione di energia (centro) e del numero di vertici risultanti dalla combinazione di coppie di tracce (destra) per b -jet (in verde), c -jet (in blu) e light-flavour jet (in rosso). Da [26].

particolarmente adatti a discriminare i b -jet da quelli di diverso flavour. La distribuzione della massa invariante non ha un picco in corrispondenza della tipica massa degli adroni contenenti b (si ricorda per il mesone B^+ , il più leggero tra essi, è di (5279.34 ± 0.12) MeV); ciò è dovuto al fatto che la massa invariante non è ricostruita utilizzando le reali masse delle particelle che provengono dal vertice secondario bensì sfruttando l'ipotesi non rigorosa di massa del pione. Inoltre, tra i prodotti di decadimento possono esservi particelle neutre, non tracciate, e nel caso di decadimenti semileptonici neutrini che sfuggono alla rivelazione.

L'algoritmo SV1 permette di calcolare altre variabili che, oltre alle tre già elencate, sono sfruttate dall'algoritmo di secondo livello descritto nella sezione 2.5.2: il numero di tracce che concorrono alla ricostruzione del vertice secondario, la distanza angolare tra l'asse del jet e la

retta che congiunge il vertice secondario al primario, la distanza tra i due vertici, la sua proiezione nel piano trasverso e la sua significanza.

Algoritmo basato sulla ricostruzione di due vertici secondari: JetFitter

L'algoritmo *JetFitter* (JF) [28] mira a ricostruire la catena di decadimenti deboli $b \rightarrow c \rightarrow s$. Esso si basa sull'assunzione, giustificata dalle loro proprietà cinematiche, che il vertice primario ed i vertici di decadimento degli adroni contenenti i quark b e c giacciono lungo la stessa traiettoria, individuata attraverso l'utilizzo di un filtro di Kalman [29]. Questa ipotesi permette di diminuire il numero di gradi di libertà del fit e rende possibile la ricostruzione della catena di decadimenti anche nello scenario in cui siano ricostruite solo una traccia associata al decadimento dell'adrone contenente b ed una associata al decadimento dell'adrone contenente c , le quali devono intersecare la stessa traiettoria.

La massa invariante delle tracce che concorrono alla ricostruzione dei vertici, la loro frazione di energia e la significanza della media tra la distanza del primo vertice dal secondo e quella del primo vertice dal terzo sono tre output utilizzati per determinare le PDF, le quali sono combinate per definire i discriminanti utilizzati per valutare le prestazioni dell'algoritmo. Analogamente al caso dell'algoritmo SV1, la massa invariante è calcolata sfruttando l'ipotesi di massa del pione per ogni particella carica. In figura 2.6 sono rappresentate le distribuzioni di tali variabili per i b -jet, i c -jet ed i light-flavour jet.

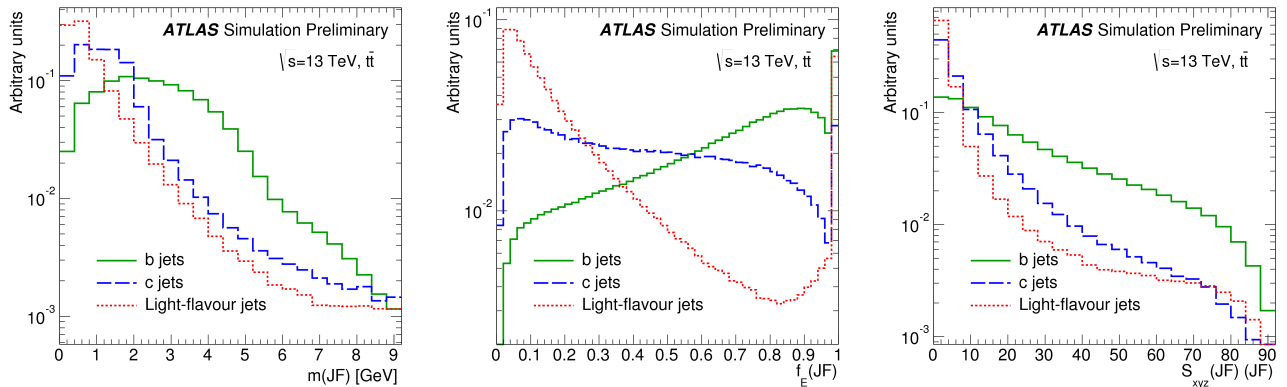


Figura 2.6: Distribuzione della massa invariante del vertice secondario ricostruito dall'algoritmo JF (sinistra), della sua frazione di energia (centro) e della significanza della media tra la distanza del primo vertice dal secondo e quella del primo vertice dal terzo (destra) per b -jet (in verde), c -jet (in blu) e light-flavour jet (in rosso). Da [26].

Altre variabili calcolate dall'algoritmo JF sono la separazione angolare tra l'asse del jet e la direzione individuata dalla somma vettoriale dei momenti delle tracce associate ai vertici secondari, il numero di tracce provenienti da vertici multi-prong, il numero di vertici individuati da una, due e più di due tracce. Tali variabili, insieme a quelle rappresentate in figura 2.6 sono utilizzate come input per l'algoritmo di secondo livello.

2.5.2 Algoritmo di secondo livello

L'algoritmo di secondo livello *MV2* [25] utilizza i risultati degli algoritmi di primo livello come input per un'analisi multivariata implementata mediante un boosted decision tree, in cui i b -jet

sono considerati come segnale e i c -jet ed i light-flavour jet come fondo. Tra le variabili utilizzate come input figurano il p_T e l' η del jet; tuttavia, le distribuzioni cinematiche del segnale e del fondo sono diverse, e ciò potrebbe essere interpretato dall'algoritmo come un discriminante tra le varie ipotesi di flavour. Per evitare tale eventualità le distribuzioni relative ai b -jet ed ai c -jet sono pesate in modo tale da riprodurre lo spettro in energia e momento dei light-flavour jet. Le altre variabili utilizzate come input sono i discriminanti degli algoritmi IP2D e IP3D e le variabili calcolate dagli algoritmi SV1 e JF elencate nei paragrafi precedenti.

La variante dell'algoritmo MV2 utilizzata durante il run 2 è denominata *MV2c20*; essa prevede che l'80% del fondo sia costituito da light-flavour jet ed il restante 20% da c -jet (da tale percentuale deriva il suffisso "c20" nel nome). Questa combinazione è il risultato di un processo di ottimizzazione in cui sono state valutate le prestazioni dell'algoritmo al variare della natura del fondo ipotizzato. L'output dell'algoritmo MV2c20 è rappresentato in figura 2.7. Su questa distribuzione vengono definiti dei tagli corrispondenti al punto di lavoro desiderato.

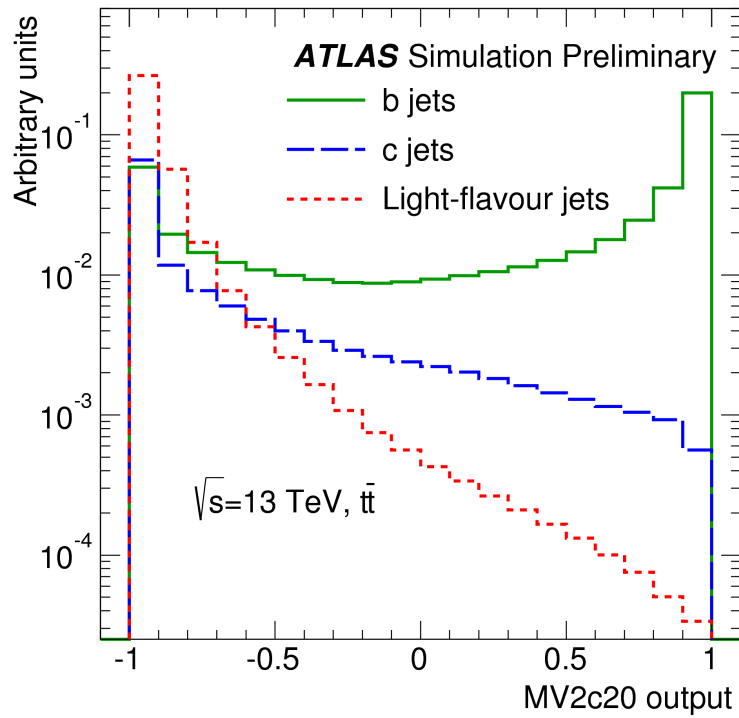


Figura 2.7: Output dell'algoritmo MV2c20 per b -jet (in verde), c -jet (in blu) e light-flavour jet (in rosso). Da [26].

2.6 Modifiche all'Inner Detector durante il run 2

L'Inner Detector, il più vicino al punto di collisione dei fasci tra i diversi rivelatori di ATLAS (si veda l'Appendice B), è adoperato per la misura delle proprietà delle particelle cariche che lo attraversano, quali la direzione, il momento e la carica, ed è pertanto immerso in un campo magnetico solenoidale di 2 T. Ha una geometria cilindrica con asse orientato nella direzione dei fasci di LHC nella quale è possibile distinguere la regione centrale, denominata *barrel*, dalle due estremità perpendicolari all'asse dei fasci, ciascuna indicata come *end-cap*. Tale geometria

permette la ricostruzione delle traiettorie di particelle con pseudorapidità $|\eta| < 2.5$. È costituito da tre diversi sottorivelatori: due dispositivi a semiconduttore, uno a pixel (il più interno) ed uno a strip di Silicio, ed un rivelatore a gas sensibile alla radiazione di transizione⁽⁴⁾ (*Transition Radiation Tracker*, TRT). I rivelatori a semiconduttore sono utilizzati per il tracciamento delle particelle cariche (maggiori informazioni sono fornite nella sezione 3.5). Il TRT sfrutta la dipendenza dell'intensità della radiazione di transizione emessa da una particella dall'inverso della sua massa per distinguere particelle cariche leggere da quelle più pesanti.

Durante il run 1, il rivelatore a pixel era costituito da 3 strati concentrici al tubo nel quale si propagano i fasci, indicato nel seguito come *beam pipe*, nella zona del barrel e da 3 strati di rivelatori in ciascun end-cap. I pixel, di geometria planare (si veda la sezione 3.5.2), avevano dimensioni $50\ \mu\text{m} \times 400\ \mu\text{m}$, con il lato più lungo nella direzione di propagazione dei fasci, ed uno spessore di $250\ \mu\text{m}$. Tra il 2013 e l'inizio del 2015, quando si è arrestata la presa dati, tale rivelatore ha subito delle modifiche volte ad adeguarlo alle condizioni in cui ATLAS ha lavorato nel run 2, conclusosi nel 2018.

Nel run 2 l'energia del centro di massa della collisione protone - protone è aumentata dagli 8 TeV raggiunti durante il run 1 a 13 TeV; in termini di luminosità istantanea, l'obiettivo era il raggiungimento di almeno $2 \times 10^{34}\ \text{cm}^{-2}\text{s}^{-1}$, ovvero il doppio della luminosità ottenuta durante il run 1. L'aumento di questi valori ha reso necessario modificare il rivelatore a pixel in modo da renderlo più resistente ai danni indotti dall'aumento di radiazione (descritti nella sezione 3.4) e da ottenere delle prestazioni almeno pari a quelle raggiunte durante il run 1 nel tracciamento delle particelle. Una prima modifica è stata lo spostamento dell'elettronica per la conversione dei segnali elettrici in segnali ottici all'esterno del rivelatore, in modo da renderla accessibile per eventuali riparazioni. Tale cambiamento è motivato dal fatto che l'aumento di luminosità e della sezione d'urto totale protone - protone determinano un incremento del flusso di particelle che può causare dei danneggiamenti dell'elettronica di lettura. Inoltre, per far fronte alla maggiore quantità di informazioni derivante dall'incremento di luminosità è stato aumentato il numero di fibre ottiche in modo da incrementare la massima frequenza di trasmissione dei dati da 80 Mbps a 160 Mbps. Una terza modifica, la più importante dal punto di vista delle prestazioni di tracciamento, è stata l'introduzione dell'*Insertable B-Layer* (IBL) [30].

L'IBL è il quarto strato cilindrico di rivelatori a pixel dell'Inner Detector di ATLAS durante il run 2, posizionato ad un raggio di 33.25 mm rispetto all'asse dei fasci, ovvero tra i tre strati già esistenti ed una nuova beam pipe di raggio interno minore rispetto al precedente; la geometria dell'Inner Detector durante il run 2 è rappresentata in figura 2.8, nella quale SCT indica il rivelatore a strip di Silicio (anche indicato come *SemiConductor Tracker*). L'IBL è costituito da 14 *staves* (doghe) in Carbonio disposte nel barrel, organizzate in modo tale da ottenere una copertura completa nell'angolo azimutale per momenti trasversi superiori a $p_T > 1\ \text{GeV}$ ed una copertura longitudinale fino a $|\eta| < 3$; ogni stave ospita 20 moduli di sensori a pixel, il loro sistema di alimentazione ed uno per il raffreddamento. In figura 2.9 è rappresentata una vista longitudinale di uno dei 14 stave. Ciascuno dei 12 moduli centrali è costituito da un rivelatore planare di spessore $200\ \mu\text{m}$ con pixel di dimensioni $50\ \mu\text{m} \times 250\ \mu\text{m}$, collegato tramite bump-bonding a due chip di lettura detti FE-I4B; i restanti 8 moduli (4 per ogni estremità) sono costituiti da rivelatori 3D⁽⁵⁾ con pixel di uguali dimensioni ma di spessore di $230\ \mu\text{m}$, ciascuno

⁽⁴⁾Radiazione emessa da una particella carica all'interfaccia di separazione tra due mezzi con diverso indice di rifrazione.

⁽⁵⁾Per la definizione di bump-bonding e di rivelatore 3D si vedano le sezioni 3.5.1 e 3.5.2.

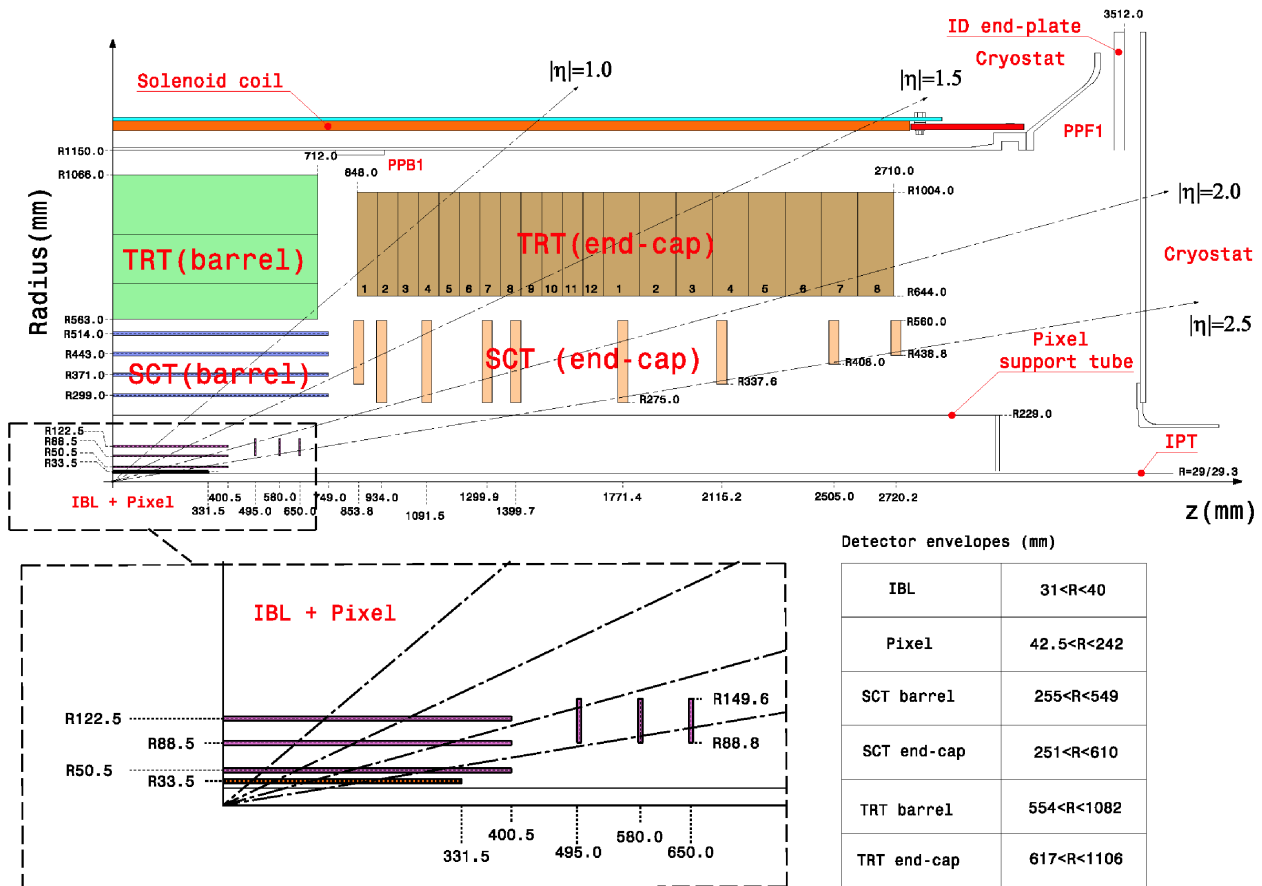


Figura 2.8: Struttura dell'Inner Detector durante il run 2. Da [30].

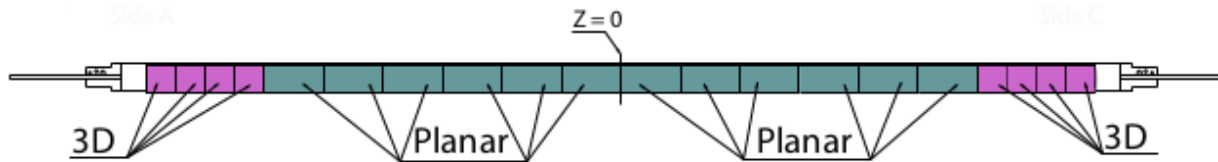


Figura 2.9: Vista longitudinale della successione di moduli di pixel alloggiati sui uno dei 14 stave di IBL. Da [30].

connesso tramite bump-bonding ad un chip FE-I4B. L'IBL è il primo esempio di applicazione della geometria 3D in rivelatori di grandi dimensioni. L'IBL è stato realizzato assemblando le varie componenti in modo che fosse indipendente dai tre strati successivi e dalla beam pipe, per permettere la sua estrazione nel caso fossero necessari interventi di manutenzione. La separazione dell'IBL dalla beam pipe è realizzata tramite l'*Inner Positioning Tube* (IPT), visibile in figura 2.8.

La geometria appena descritta ed i materiali utilizzati per la sua realizzazione rispondono all'esigenza di minimizzare il numero di lunghezze di radiazione percorse dalle particelle durante la loro propagazione all'interno del rivelatore; ciò è fondamentale per limitare l'impatto negativo che l'interazione con il materiale del rivelatore ha sul loro tracciamento per effetto dello scattering multiplo coulombiano e, in misura minore, della progressiva perdita di energia. Indicando con

X_0 la lunghezza di radiazione, la distanza percorsa all'interno dell'IBL da una particella prodotta alla coordinata $z = 0$ con $\eta = 0$ corrisponde a $1.88\% X_0$; ciò rappresenta una riduzione del 30% rispetto alla quantità di materiale corrispondente ad uno degli strati successivi dell'Inner Detector. Tale risultato è stato raggiunto grazie all'utilizzo della nanoschiama di Carbonio⁽⁶⁾ per la realizzazione dei 14 supporti ed all'introduzione di un sistema di raffreddamento basato su CO_2 , che richiede l'utilizzo di una minore quantità di gas rispetto al sistema basato su C_3F_8 utilizzato per gli altri strati del rivelatore a pixel.

2.6.1 Impatto sulla ricostruzione delle tracce

L'introduzione dell'IBL ha apportato dei miglioramenti alla ricostruzione delle tracce.

Un primo progresso deriva dall'avere a disposizione un nuovo punto attraverso cui estrapolare la traccia di una particella in prossimità del punto di collisione dei fasci; ciò è particolarmente importante a bassi valori di momento trasverso, laddove la ricostruzione delle tracce è maggiormente affetta dai problemi derivanti dallo scattering multiplo delle particelle sul materiale del rivelatore. L'aggiunta di un nuovo strato di rivelatori a pixel determina dunque un miglioramento della risoluzione dei parametri di impatto trasversi e longitudinali, fondamentali per gli algoritmi di b -tagging (come spiegato in sezione 2.5). Tale risoluzione, inoltre, è fortemente influenzata dalla granularità del rivelatore, dettata dalla distanza tra i centri di due pixel successivi, detta *pitch*; pertanto, l'utilizzo di pixel di dimensioni $50 \mu\text{m} \times 250 \mu\text{m}$ comporta un ulteriore miglioramento sulla risoluzione del parametro di impatto longitudinale per ogni valore di p_T , dal momento che le dimensioni dei pixel degli strati successivi dell'Inner Detector sono pari a $50 \mu\text{m} \times 400 \mu\text{m}$. In figura 2.10 è rappresentato un confronto tra la risoluzione sul parametro di impatto trasverso (sinistra) e longitudinale (destra) in funzione del p_T della traccia relativamente ai dati raccolti durante il run 1 (in nero) ed il run 2 (in rosso). L'andamento osservato può

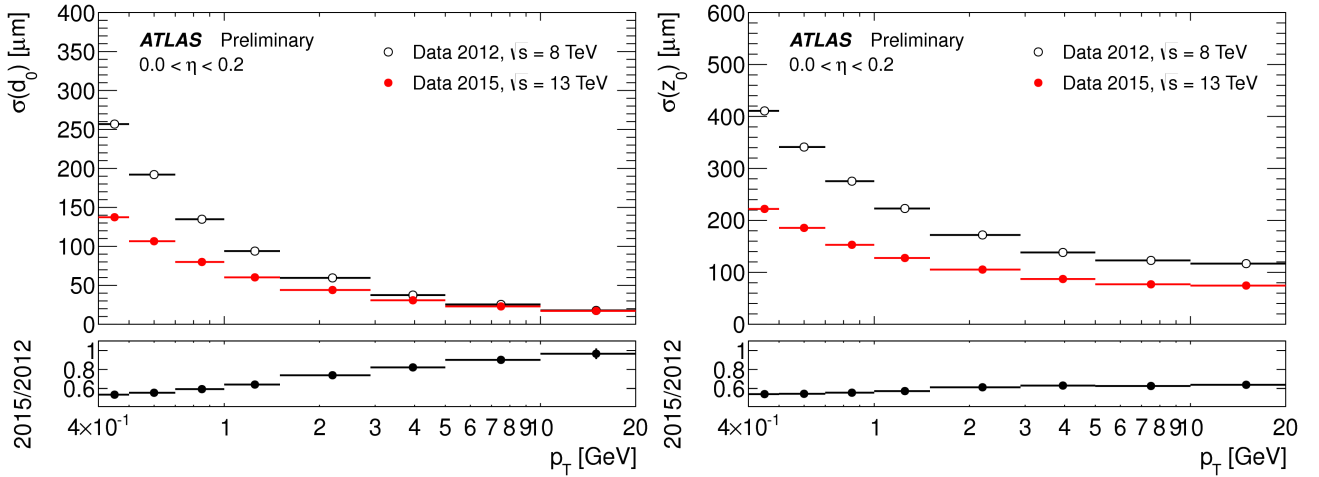


Figura 2.10: Confronto tra la risoluzione sul parametro di impatto trasverso (sinistra) e longitudinale (destra) in funzione del p_T della traccia nel run 1 (in nero) e quella nel run 2 (in rosso). Nei pannelli inferiori è rappresentato il rapporto tra le distribuzioni nei relativi pannelli superiori. Da [31].

essere descritto dalla relazione $A \oplus B/p_T$, in cui il termine A rappresenta la risoluzione intrinseca del rivelatore, determinata dal pitch, mentre B descrive gli effetti dello scattering multiplo nel

⁽⁶⁾Allotropo del Carbonio di bassa densità.

materiale del rivelatore, dominanti a basso momento trasverso. Il termine A domina invece ad alti valori di p_T ; il fatto che in tale regione il miglioramento indotto dall'IBL sia maggiore per il parametro di impatto longitudinale è dovuto alla diminuzione del pitch in tale direzione.

In figura 2.11 il confronto tra le risoluzioni durante i due run è mostrato in funzione della pseudorapidità della traccia. Anche in questi grafici è evidente il complessivo miglioramento

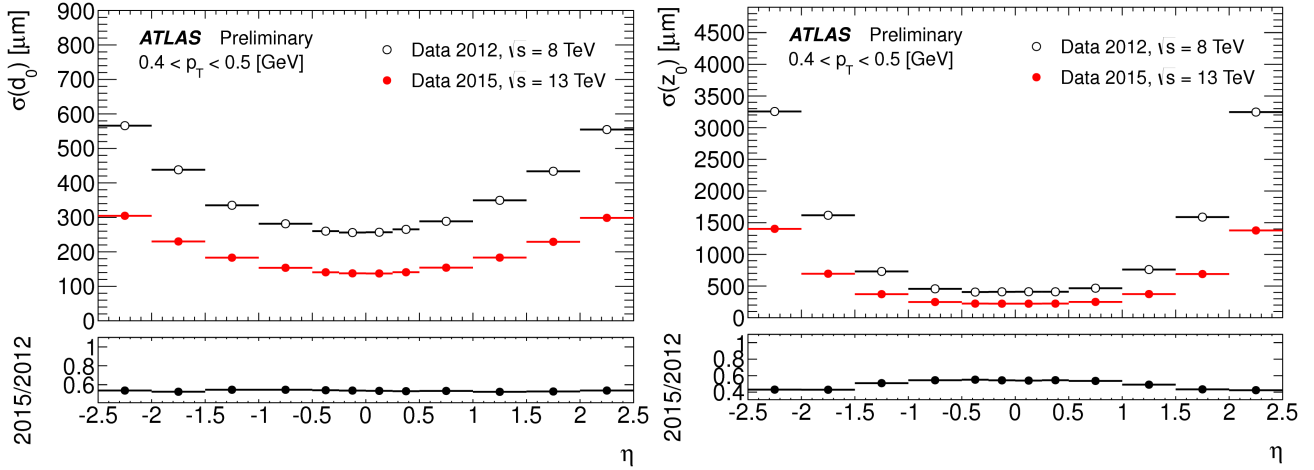


Figura 2.11: Confronto tra la risoluzione sul parametro di impatto trasverso (sinistra) e longitudinale (destra) in funzione dell' $|\eta|$ della traccia nel run 1 (in nero) e nel run 2 (in rosso). Nei pannelli inferiori è rappresentato il rapporto tra le distribuzioni nei relativi pannelli superiori. Da [31].

indotto dall'introduzione dell'IBL il quale tuttavia, a differenza del caso precedente, è pressoché costante nell'intervallo considerato. Inoltre, si osserva come la risoluzione peggiori ad alti valori di $|\eta|$ in entrambi i run. Ciò è dovuto al fatto che le tracce con valori di $|\eta|$ più piccoli sono estrapolate da misure più vicine al vertice della collisione, come visibile nell'ingrandimento del rivelatore a pixel in figura 2.8, e pertanto la misura del punto di massimo avvicinamento al vertice è più precisa.

2.6.2 Impatto sull'efficienza di b -tagging

Il miglioramento della risoluzione del parametro di impatto indotto dall'IBL descritto nella sezione 2.6 si traduce in un miglioramento del b -tagging. Nelle figure 2.12, 2.13, 2.14 sono rappresentate le prestazioni su dati simulati degli algoritmi di b -tagging nella geometria del rivelatore del run 2 (in blu), in confronto con i risultati relativi alla geometria precedente all'installazione dell'IBL (in rosso). I grafici illustrati si riferiscono all'algoritmo di secondo livello MV2c20 applicato alla simulazione di eventi $t\bar{t}$, in cui si è tenuto conto delle differenze di energia del centro di massa e geometria dell'Inner Detector tra i due periodi di presa dati.

In figura 2.12 è rappresentata la reiezione di light-flavour jet (a sinistra) e c -jet (a destra) in funzione dell'efficienza di b -tagging nella geometria del run 1 (in rosso) ed in quella del run 2 (in blu). Risulta evidente come l'introduzione dell'IBL comporti un miglioramento nella reiezione di light-flavour jet di un fattore compreso tra 1 e 5 nell'intervallo di efficienza di b -tagging considerato. Anche la reiezione di c -jet beneficia dall'introduzione dell'IBL, ma in misura minore. Ciò è dovuto al fatto che la segnatura di un c -jet è simile a quella di un b -jet (anche gli adroni contenenti c hanno, infatti, elevate lunghezze di decadimento e parametri di impatto) e pertanto nella distinzione di un b -jet da un c -jet giocano un ruolo importante le logiche degli

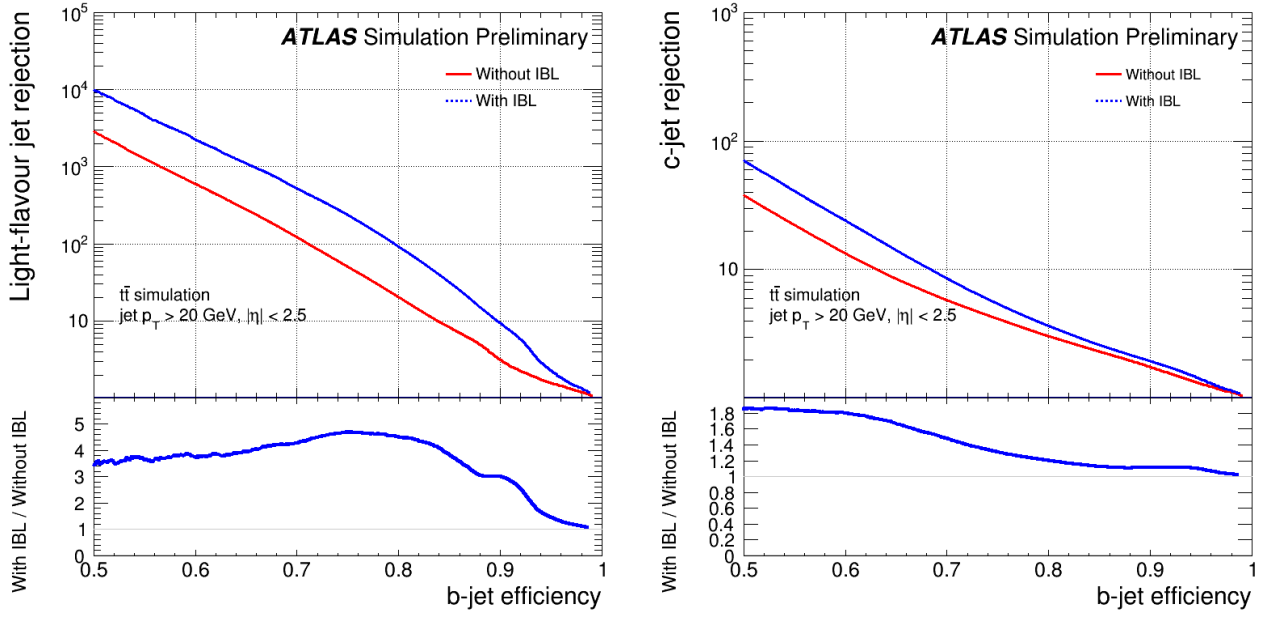


Figura 2.12: Confronto tra la reiezione di light-flavour jet (sinistra) e c-jet (destra) in funzione dell'efficienza di b-tagging durante il run 1 (in rosso) ed il run 2 (in blu). Da [32].

algoritmi utilizzati oltre alla risoluzione; di conseguenza, il rapporto tra la reiezione ottenuta grazie all'IBL e quella precedente alla sua installazione non supera 1.8.

In figura 2.13 è rappresentato un confronto tra la reiezione di light-flavour jet (a sinistra) e di c-jet (a destra) in funzione del momento trasverso del jet nelle due diverse configurazioni geometriche per un punto di lavoro dell'algoritmo di secondo livello tale che l'efficienza di b-tagging sia del 70%. Si osserva un netto miglioramento indotto dall'IBL per bassi valori del

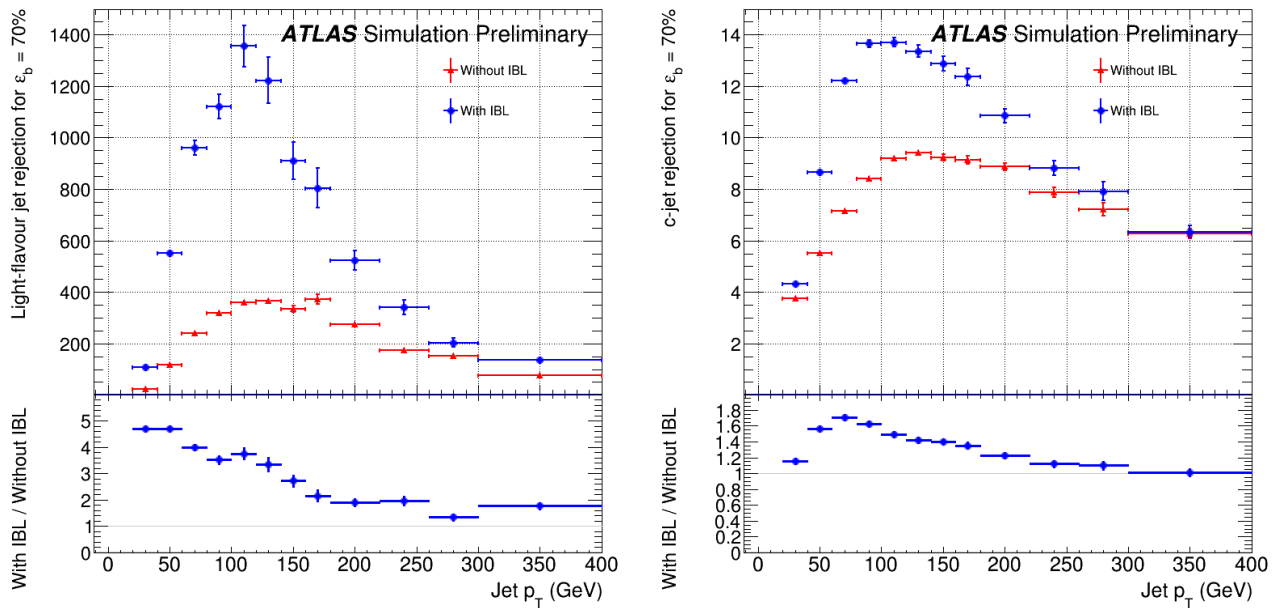


Figura 2.13: Confronto tra la reiezione di light-flavour jet (sinistra) e c-jet (destra) in funzione del momento trasverso del jet durante il run 1 (in rosso) ed il run 2 (in blu). Le distribuzioni sono relative ad un'efficienza di b-tagging del 70%. Da [32].

momento trasverso, il quale riflette l'andamento della risoluzione osservato in figura 2.10.

In figura 2.14 la reiezione di light-flavour jet (sinistra) e c -jet (destra) è rappresentata in funzione dell' $|\eta|$ del jet. Il miglioramento introdotto dall'IBL in queste misure è pressoché

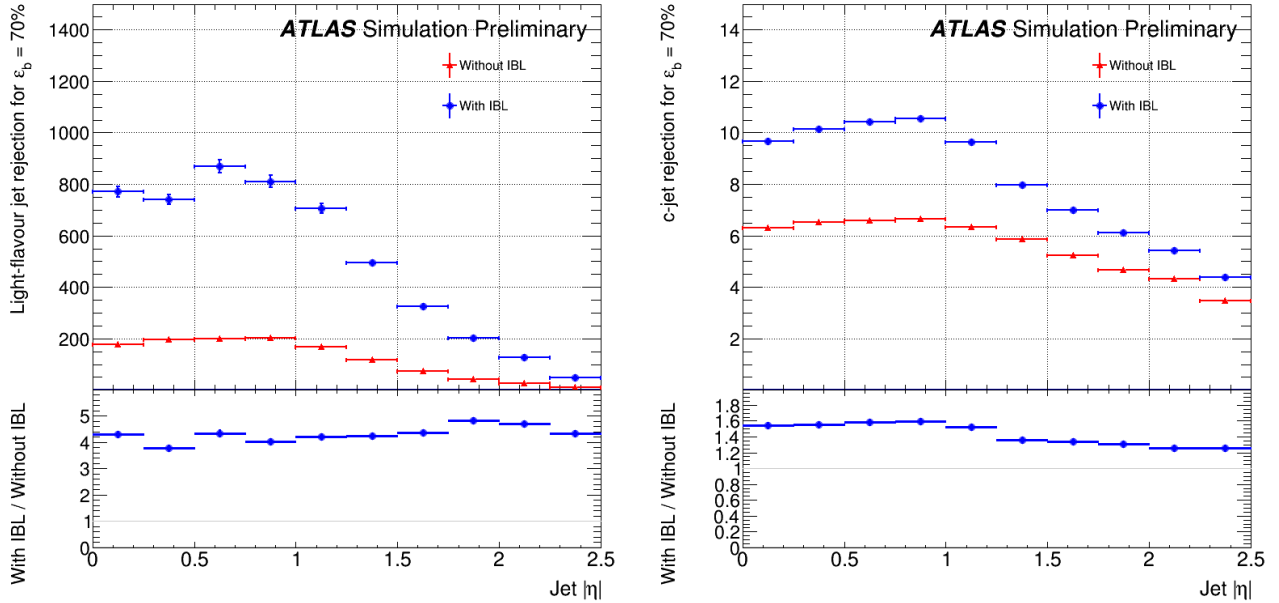


Figura 2.14: Confronto tra la reiezione di light-flavour jet (sinistra) e c -jet (destra) in funzione della pseudorapidità del jet durante il run 1 (in rosso) ed il run 2 (in blu). Le distribuzioni sono relative ad un'efficienza di b -tagging del 70%. Da [32].

costante al variare della pseudorapidità; ciò è coerente con quanto già osservato in figura 2.11.

2.7 Inner Tracker per HL–LHC

Come accennato nella sezione 1.1.2, la collisione tra i fasci di protoni ad LHC avverrà a partire dal 2027 ad una luminosità istantanea maggiore rispetto a quella prevista per il run 3; essa raggiungerà il valore di $\mathcal{L} = 7.5 \times 10^{34} \text{ cm}^{-2}\text{s}^{-1}$, che equivale ad un incremento di un fattore pari a circa 3.5 rispetto alla precedente. L'obiettivo è quello di raccogliere una quantità di dati corrispondente ad una luminosità integrata $L = 4000 \text{ fb}^{-1}$ in circa 12 anni di presa dati. Il run caratterizzato da tali valori di luminosità prende il nome di *High-Luminosity LHC* (HL–LHC); l'energia nel centro di massa sarà uguale a quella che sarà già raggiunta durante il run 3, pari a 14 TeV.

L'aumento di luminosità ad HL–LHC comporta un ulteriore incremento del pile-up di un fattore 10 rispetto a quello del run 2: si prevede che il numero medio di interazioni in corrispondenza di ogni collisione tra i fasci di protoni sarà pari a 200. La figura 2.15 mostra il confronto tra un evento tipico registrato nel run 1 (a sinistra), nel quale il pile-up era inferiore a quello del run 2, e la simulazione di un evento atteso ad HL–LHC (a destra). È evidente come l'aumento del numero di tracce in ogni collisione renda necessario adeguare i vari rivelatori.

L'Inner Detector di ATLAS descritto nella sezione 2.6 sarà sostituito da una nuovo rivelatore, realizzato interamente in Silicio, detto *Inner Tracker* (ITk). Esso è stato progettato in modo tale da ottenere delle prestazioni di tracciamento almeno pari a quelle raggiunte durante il run 2

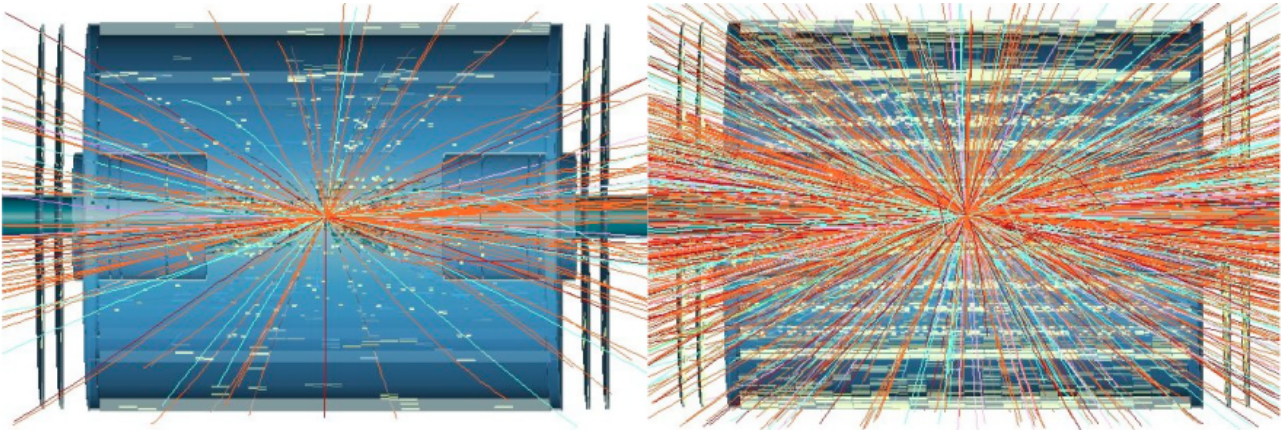


Figura 2.15: Confronto tra un tipico evento registrato nel run 1 (sinistra) e la simulazione di un evento atteso ad HL-LHC (destra).

nonostante l'aumento di pile-up ed ammettendo la possibilità di difetti di alcune componenti, causati dall'elevata dose di radiazione prevista.

La geometria dell'ITk, anch'essa cilindrica, è rappresentata a sinistra in figura 2.16, nella quale è illustrato un singolo quadrante; l'origine rappresenta il punto di interazione protone - protone, l'asse orizzontale è orientato lungo la linea dei fasci, l'asse verticale verso l'alto. ITk è costituito

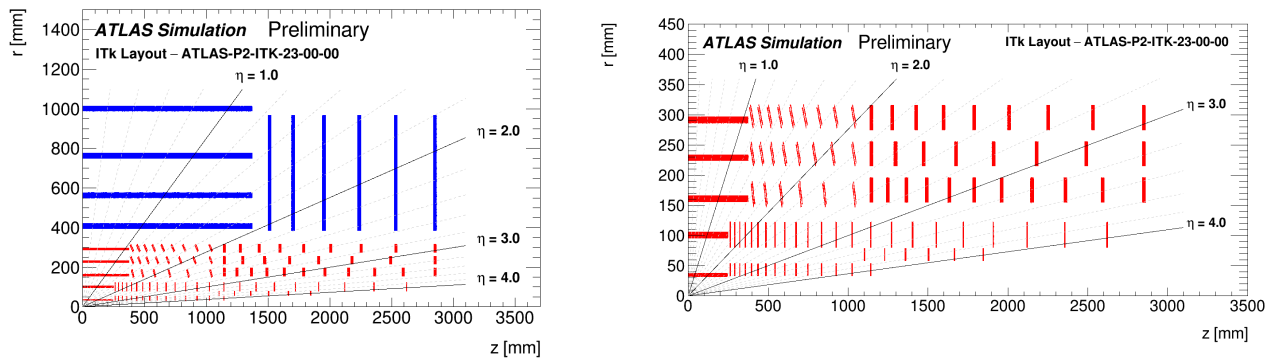


Figura 2.16: Sinistra: geometria dell'Inner Tracker Detector progettato per HL-LHC, in cui sono visibili il rivelatore a strip (in blu) e quello a pixel (in rosso). Destra: dettaglio del rivelatore a pixel. In entrambi i casi è rappresentato solo un quadrante; l'origine rappresenta il punto di interazione protone - protone, l'asse orizzontale è orientato lungo la linea dei fasci, l'asse verticale indica il raggio misurato a partire dal punto di collisione. Da [33].

da due sottorivelatori, entrambi presenti sia nel barrel che negli end-cap: il rivelatore a strip, indicato in blu in figura, con copertura fino a $|\eta| < 2.7$, ed il rivelatore a pixel, indicato in rosso, con copertura fino a $|\eta| < 4.0$. I due sottorivelatori sono divisi dal *Pixel Support Tube* (PST), un cilindro realizzato con fibre di carbonio. I due strati più interni del rivelatore a pixel sono separati dagli altri dall'*Inner Support Tube* (IST) in modo tale da consentire la loro estrazione; ciò risponde all'esigenza di sostituirli dopo aver raccolto una statistica corrispondente a 2000 fb^{-1} come conseguenza del danneggiamento causato dall'elevato livello di radiazione al quale saranno stati sottoposti.

A destra in figura 2.16 è rappresentato un ingrandimento della geometria del rivelatore a pixel, in cui è evidente la presenza di moduli inclinati tra il barrel e l'end-cap; questi rendono

più graduale la transizione tra le due regioni limitando l'effetto negativo che la loro separazione ha sulla ricostruzione delle tracce. La geometria descritta risponde all'obiettivo di ottenere una copertura completa per particelle che provengono da un beam spot con estensione in z tra -15 cm e $+15$ cm.

In figura 2.17 è rappresentato il numero di segnali registrati dai rivelatori a pixel e a strip in funzione di $|\eta|$. I risultati sono relativi alla simulazione di singoli muoni con $p_T = 15$ GeV

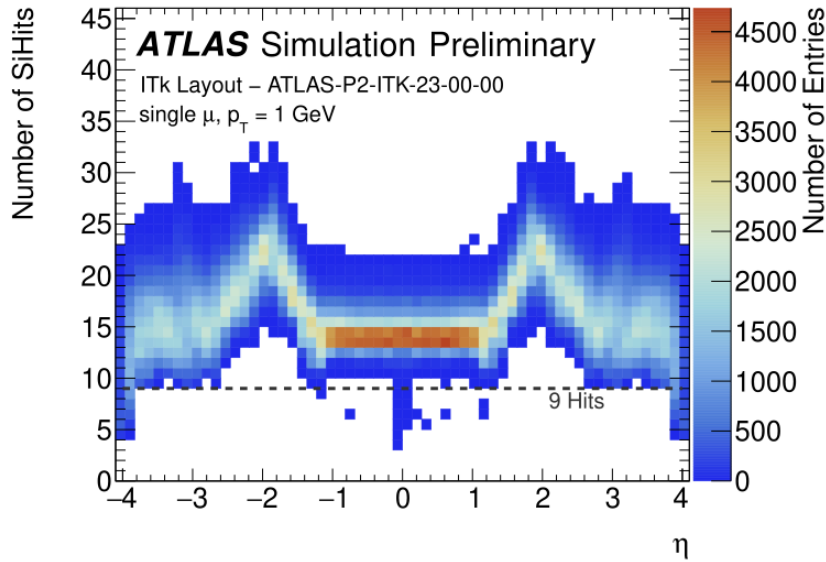


Figura 2.17: Numero di hit registrati da ITk relativamente alla simulazione di singoli muoni con $p_T = 15$ GeV prodotti con una distribuzione uniforme tra 0 mm e 2 mm in r e con $z = -15$ cm, 0 cm, 15 cm. Da [33].

prodotti con una distribuzione uniforme tra 0 mm e 2 mm in r e con $z = -15$ cm, 0 cm, 15 cm. Affinché si possa ricostruire la traccia di una particella è necessaria l'osservazione di almeno 9 hit; si osserva che tale richiesta è soddisfatta in quasi tutto l'intervallo $|\eta| < 4$, con piccole eccezioni dovute al passaggio delle particelle nelle zone di separazione tra i vari moduli.

2.7.1 Impatto sull'efficienza di b -tagging

I risultati presentati in questo paragrafo sono basati su una geometria del rivelatore leggermente diversa da quella finale, che è stata illustrata in figura 2.16. La principale differenza tra le due è rappresentata dalla riduzione della distanza dello strato più interno di rivelatori a pixel dal beam spot: nel progetto finale è di 34 mm, nella proposta precedente di 39 mm. Tuttavia, la geometria definitiva è stata approvata nel gennaio 2020, e non sono ancora stati resi pubblici gli studi sulle prestazioni del b -tagging aggiornati alla versione finale del progetto. Nei grafici illustrati nel seguito si segneranno, ove possibile, le differenze previste a causa del cambiamento della geometria.

In figura 2.18 è rappresentato il numero di lunghezze di interazione nucleare percorse da una particella prima di soddisfare i criteri sul minimo numero di hit necessario alla ricostruzione della sua traccia in funzione della pseudorapidità. La distribuzione relativa ad ITk (in verde) è confrontata con quella relativa all'Inner Detector del run 2 (in blu); nel pannello inferiore è rappresentato il rapporto tra le due. È evidente come, oltre ad aumentare la copertura in $|\eta|$, in

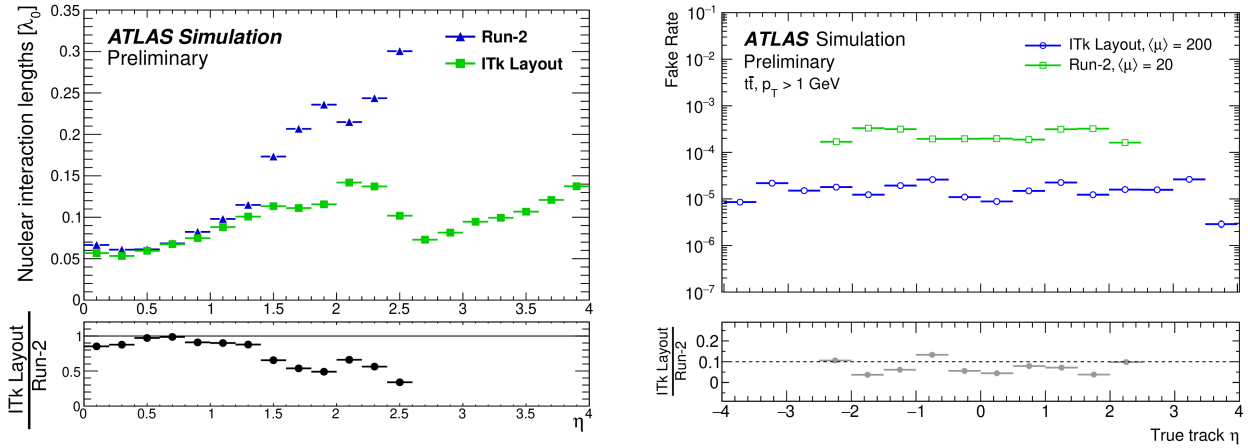


Figura 2.18: Sinistra: confronto tra il numero di lunghezze di interazione percorse da una particella prima di soddisfare i criteri sul minimo numero di hit necessario alla ricostruzione della sua traccia in funzione della pseudorapidità nel run 2 (in blu) e in HL-LHC (in verde). Destra: confronto tra la fake rate in funzione della pseudorapidità della traccia vera nel run 2 (in verde) e ad HL-LHC (in blu). Nei pannelli inferiori sono rappresentati i rapporti tra le distribuzioni nei relativi pannelli superiori. Da [34].

ITk diminuisca il numero di lunghezze di interazione percorse, con una conseguente diminuzione degli effetti di interazione delle particelle con il materiale del rivelatore, i quali compromettono la loro ricostruzione. Tale diminuzione è il risultato dell'ottimizzazione della disposizione dei rivelatori e di ulteriori accorgimenti. Tra questi riveste particolare importanza l'introduzione di un sistema di alimentazione seriale dei chip di lettura del rivelatore a pixel che consente di diminuire il numero di cavi utilizzati (tale sistema è discusso nella sezione 4.3).

Un possibile problema determinato dall'aumento di pile-up previsto per HL-LHC è l'aumento della *fake rate*, ovvero della frequenza di associazione di tracce fittizie all'evento considerato. Tuttavia, la richiesta di osservazione di un elevato numero di hit registrati in rivelatori di alta precisione, quali quelli al Silicio, permette di far fronte a tale problema, come dimostrato in figura 2.18 (destra). In tale simulazione la fake rate è definita come la frazione di tracce ricostruite ma non corrispondenti ad una particella effettivamente simulata all'interno del detector; è evidente come essa sia inferiore in ITk (in blu) rispetto a quella ottenuta con l'Inner Detector del run 2 (in verde), nonostante il pile-up medio considerato sia di 200 e 20 rispettivamente.

In figura 2.19 è rappresentata la dipendenza della risoluzione sul parametro di impatto trasverso dall' $|\eta|$ della traccia. Il grafico a sinistra è relativo alla simulazione di singoli muoni di momento trasverso pari a 1 GeV mentre nel plot a destra il loro p_T è 100 GeV. Le tre distribuzioni visibili in entrambe le figure sono relative alla geometria del run 2 (in nero) e a quella di ITk in cui i pixel hanno dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$ (in verde) oppure $25 \mu\text{m} \times 100 \mu\text{m}$ (in blu). Le ultime due fanno riferimento a due diverse ipotesi avanzate per la realizzazione del rivelatore a pixel. Per bassi momenti trasversi, dove l'interazione con il materiale ha effetti più importanti, le prestazioni del run 2 sono migliori rispetto a quelle previste per ITk; ciò è dovuto al fatto che l'IBL è collocato ad una distanza di circa 33 mm dall'asse dei fasci, mentre lo strato più interno nella geometria di ITk considerata per questa simulazione si trova a 39 mm. La scelta di diminuire tale distanza nella geometria finale risponde all'esigenza di eguagliare le prestazioni ottenute nel run 2. Per alti valori di momenti trasversi la risoluzione ottenuta considerando pixel di $50 \mu\text{m} \times 50 \mu\text{m}$ in ITk è confrontabile con quella relativa al run 2, mentre pixel di dimensioni

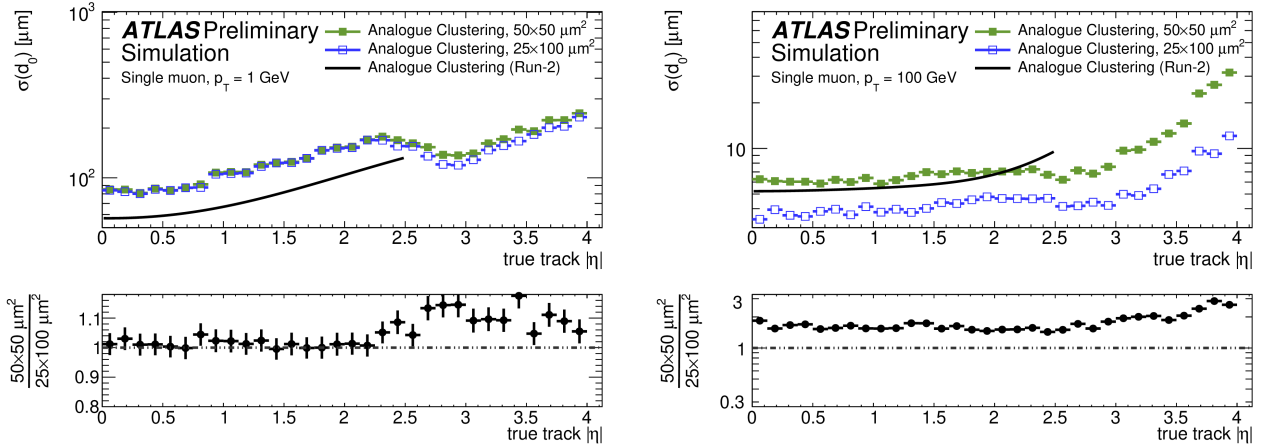


Figura 2.19: Dipendenza della risoluzione sul parametro di impatto trasverso dalla pseudorapidità della traccia nel caso della simulazione di singoli muoni con $p_T = 1$ GeV (sinistra) oppure $p_T = 100$ GeV (destra). Le distribuzioni sono relative al run 2 (in nero), ad ITk con pixel di dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$ (in verde) oppure $25 \mu\text{m} \times 100 \mu\text{m}$ (in blu). Nei pannelli inferiori sono rappresentati i rapporti tra le due distribuzioni relative ad ITk nei relativi pannelli superiori. Da [34].

$25 \mu\text{m} \times 100 \mu\text{m}$ permettono di ottenere un miglioramento di un fattore 2 rispetto alle altre due configurazioni; ciò è dovuto alla migliore risoluzione nella coordinata trasversa di tali pixel (si ricorda che le dimensioni dei pixel nell'IBL sono di $50 \mu\text{m} \times 250 \mu\text{m}$).

In figura 2.20 è rappresentato un confronto tra la risoluzione del parametro di impatto longitudinale in funzione di $|\eta|$ per le tre geometrie considerate. Analogamente alla figura

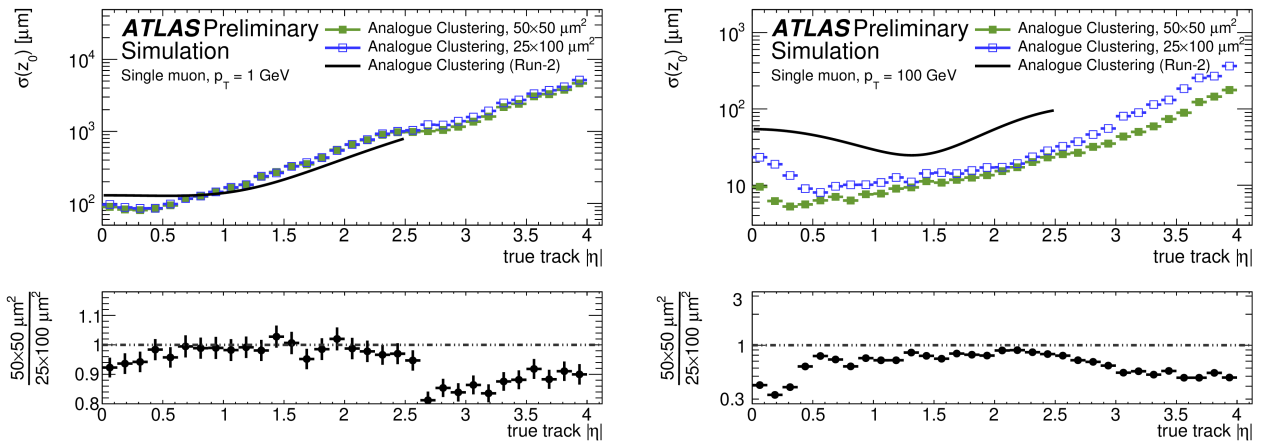


Figura 2.20: Dipendenza della risoluzione sul parametro di impatto longitudinale dalla pseudorapidità della traccia nel caso di singoli muoni con $p_T = 1$ GeV (sinistra) oppure $p_T = 100$ GeV (destra). Le distribuzioni sono relative al run 2 (in nero), ad ITk con pixel di dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$ (in verde) oppure $25 \mu\text{m} \times 100 \mu\text{m}$ (in blu). Nei pannelli inferiori sono rappresentati i rapporti tra le due distribuzioni relative ad ITk. Da [34].

precedente, il grafico a sinistra fa riferimento alla simulazione di singoli muoni di momento trasverso pari a 1 GeV mentre quello a destra è relativo a singoli muoni di $p_T = 100$ GeV. Il netto miglioramento visibile in ITk per muoni di momento trasverso pari a 100 GeV è attribuibile

alla diminuzione della dimensione dei pixel nella direzione longitudinale per entrambe le ipotesi considerate.

Il fatto che ITk sia realizzato interamente in Silicio risponde all'esigenza di migliorare la misura del momento delle particelle che lo attraversano rispetto a ciò che accade nell'Inner Detector del run 2, in cui lo strato più esterno è rappresentato dal Transition Radiation Tracker. Il risultato previsto è confermato dai grafici in figura 2.21, in cui è rappresentata la dipendenza della risoluzione relativa sull'inverso del momento trasverso (q indica la carica elettrica della particella, determinata dal verso in cui si avvolge l'elica da essa descritta nel campo magnetico, noto il suo verso) in funzione di $|\eta|$ per singoli muoni di $p_T = 1$ GeV (sinistra) oppure $p_T = 100$ GeV. Le dimensioni dei pixel giocano un ruolo importante ad alti valori di p_T , laddove la risoluzione è

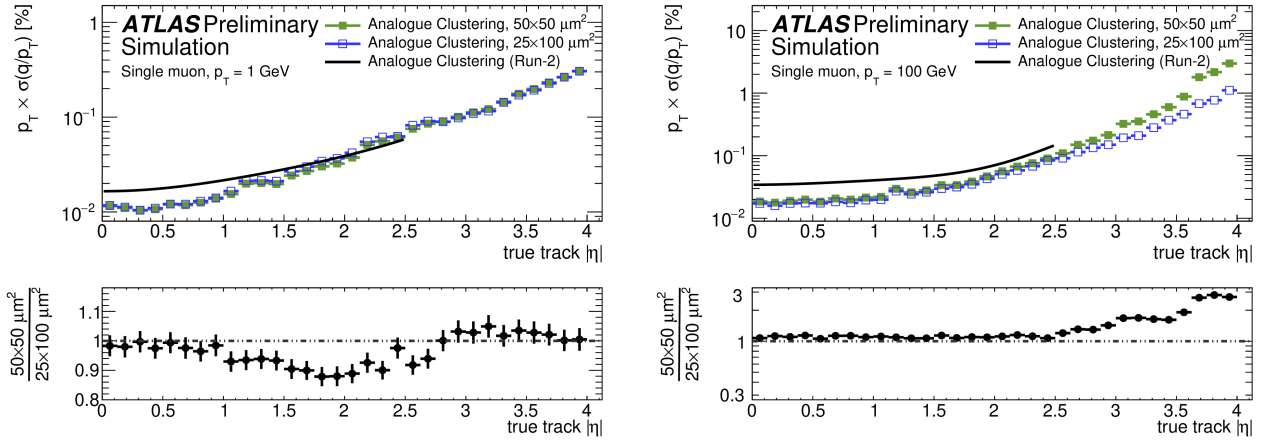


Figura 2.21: Dipendenza della risoluzione relativa sul momento trasverso di singoli muoni con $p_T = 1$ GeV (sinistra) oppure $p_T = 100$ GeV (destra) dalla pseudorapidità della traccia. Le distribuzioni sono relative al run 2 (in nero), ad ITk con pixel di dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$ (in verde) oppure $25 \mu\text{m} \times 100 \mu\text{m}$ (in blu). Nei pannelli inferiori sono rappresentati i rapporti tra le due distribuzioni relative ad ITk. Da [34].

maggiormente determinata da quella intrinseca del detector. A bassi valori di momento trasverso, invece, predominano gli effetti di interazione con il materiale del rivelatore e la differenza tra pixel di diverse dimensioni è meno evidente.

I risultati rappresentati nelle figure precedenti influenzano le prestazioni di flavour tagging. In figura 2.22 (sinistra) è rappresentata la reiezione di light-flavour jet in funzione dell'efficienza di b -tagging per diversi range di $|\eta|$: $0 < |\eta| < 1$ (in nero), $1 < |\eta| < 2$ (in blu), $2 < |\eta| < 3$ (in rosso), $3 < |\eta| < 4$ (in verde). Le linee continue sono relative a pixel di dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$, quelle tratteggiate a pixel di dimensioni $25 \mu\text{m} \times 100 \mu\text{m}$. I risultati ottenuti si riferiscono alla simulazione di eventi $t\bar{t}$ con un pile-up medio di 200 ed all'utilizzo dell'algoritmo di b -tagging IP3D. È evidente come la reiezione di light-flavour jet peggiori all'aumentare del modulo della pseudorapidità; ciò è attribuibile al peggioramento della risoluzione sul parametro di impatto ad alti $|\eta|$ osservata nelle figure 2.19 e 2.20. Inoltre, la reiezione ottenuta considerando pixel di dimensione $25 \mu\text{m} \times 100 \mu\text{m}$ è sempre leggermente superiore a quella relativa a pixel di dimensione $50 \mu\text{m} \times 50 \mu\text{m}$; per questo motivo si è scelto di confrontare i risultati relativi a tale geometria con la reiezione ottenuta durante il run 2 per efficienze di ricostruzione di b -jet pari a 60%, 70%, 77% e 85%, ovvero ai punti di lavoro considerati durante il run 2. Tale confronto è rappresentato a destra in figura 2.22, in cui i risultati relativi al run 2 sono indicati con croci verdi. Alle curve raffigurate a sinistra è stata aggiunta con un tratteggio nero quella relativa

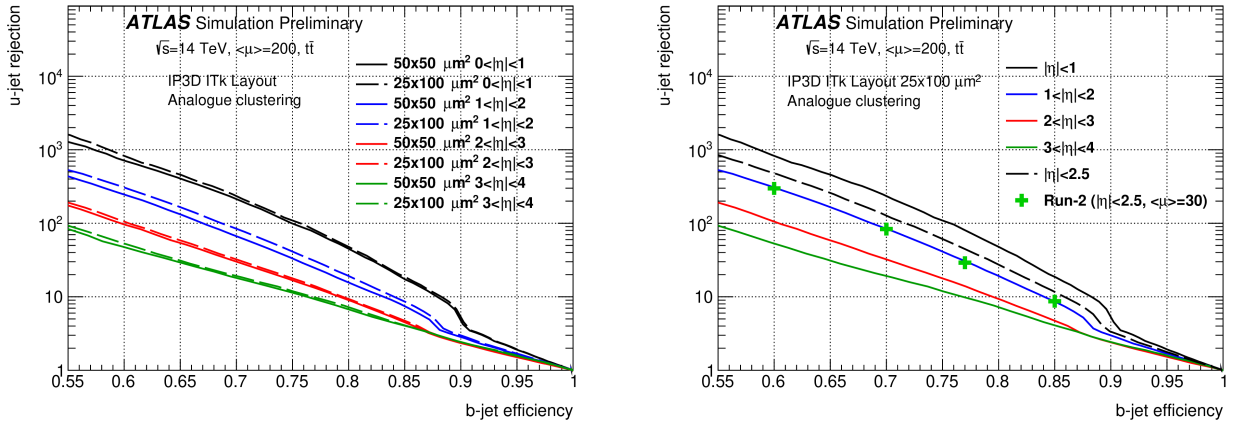


Figura 2.22: Sinistra: reiezione di light-flavour jet in funzione dell'efficienza di b -tagging per diversi valori di pseudorapidità e di dimensioni dei pixel di ITk. Destra: confronto tra la reiezione di light-flavour jet in funzione dell'efficienza di b -tagging ottenuta pixel di dimensioni $25 \mu\text{m} \times 100 \mu\text{m}$ in ITk e con l'Inner Detector. Da [34].

a $|\eta| < 2.5$, corrispondente alla copertura dell'Inner Detector, in modo da poter effettuare il confronto con il run precedente; è evidente come tale reiezione sia migliore rispetto a quella ottenuta durante il run 2 per ogni punto di lavoro considerato.

Sulla base dei grafici illustrati in questa sezione è possibile concludere che la geometria $25 \mu\text{m} \times 100 \mu\text{m}$ consente di ottenere migliori risoluzioni sul parametro di impatto trasverso ed efficienze di b -tagging. D'altra parte, pixel di dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$ forniscono migliori risoluzioni sul parametro di impatto longitudinale; essi, inoltre, permettono di sopprimere in maniera più efficiente il pile-up. La prima geometria sarà utilizzata nel primo strato del barrel, la seconda negli strati successivi e negli end-cap come risultato degli studi di ottimizzazione del progetto del tracciatore di ATLAS per HL-LHC conclusi tra la fine del 2019 e l'inizio del 2020.

2.8 Impatto del b -tagging nella ricerca di coppie di Higgs ad HL–LHC

Sulla base delle caratteristiche di progetto dell'ITk è previsto un aumento dell'efficienza di b -tagging dell'8% nelle condizioni di run di HL–LHC rispetto a quella ottenuta utilizzando l'Inner Detector [35]; considerando che i quark b sono presenti nello stato finale dei principali canali di decadimento delle coppie di Higgs ricercate è ragionevole assumere che tale miglioramento influenzerà i risultati di analisi analoghe a quelle descritte nel capitolo 1. ATLAS ha condotto degli studi volti ad estrapolare i risultati relativi alla produzione tramite gluon-gluon fusion alla luminosità di HL–LHC, assumendo un pile-up medio per ogni interazione dei fasci di 200, il miglioramento sull'efficienza di b -tagging ed un campione di dati corrispondente a 3000 fb^{-1} . I canali considerati sono $HH \rightarrow b\bar{b}b\bar{b}$, $HH \rightarrow b\bar{b}W^+W^-$ e $HH \rightarrow b\bar{b}\tau^+\tau^-$ [36].

In figura 2.23 è rappresentato l'opposto del logaritmo naturale del rapporto tra il valore di massima verosimiglianza per un dato k_λ e quello relativo all'ipotesi $k_\lambda = 1$ in funzione di k_λ . Il grafico a sinistra è relativo al caso in cui si considerino solo le incertezze statistiche e quello a destra al caso in cui si includano anche le incertezze sistematiche. Le linee tratteggiate indicano

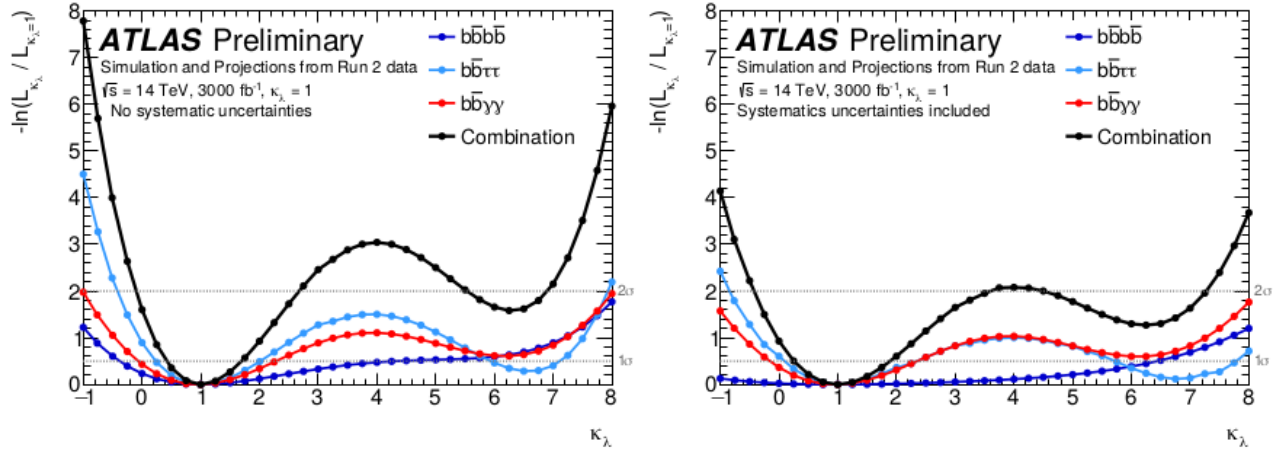


Figura 2.23: Opposto del logaritmo naturale del rapporto tra il valore di massima verosimiglianza per un dato k_λ e quello relativo al caso $k_\lambda = 1$ in funzione di k_λ nel caso in cui si considerino solo le incertezze statistiche (sinistra) o anche quelle sistematiche (destra). Le linee tratteggiate indicano i valori relativi agli estremi degli intervalli di confidenza ad 1σ e 2σ . Da [36].

i valori relativi agli estremi degli intervalli di confidenza ad 1σ e 2σ . A partire da tali grafici si estrapolano i possibili intervalli di appartenenza di k_λ riportati in tabella 2.1.

Inceteezze	1σ CL	2σ CL
Statistiche	$0.4 \leq k_\lambda \leq 1.7$	$-0.10 \leq k_\lambda \leq 2.7 \cup 5.5 \leq k_\lambda \leq 6.9$
Statistiche + sistematiche	$0.25 \leq k_\lambda \leq 1.9$	$-0.4 \leq k_\lambda \leq 7.3$

Tabella 2.1: Estrapolazione degli intervalli di confidenza ad 1σ e a 2σ per il parametro k_λ alla luminosità di HL–LHC.

In figura 2.24 è riportata la significanza con cui ci si aspetta di osservare coppie di bosoni di Higgs in funzione del valore di k_λ . Le linee tratteggiate rappresentano le soglie di 3σ e 5σ ed analogamente al caso precedente il grafico a sinistra è relativo al caso in cui si considerino solo le incertezze statistiche e quello a destra al caso in cui si includano anche quelle sistematiche. È possibile osservare che, in assenza di evoluzioni significative della strategia di analisi e di controllo delle incertezze sistematiche, solo una combinazione dei risultati di ATLAS e CMS potrà garantire di accertare un valore di k_λ consistente con la predizione del Modello Standard.

2.9 Evoluzione futura degli algoritmi di b -tagging

Gli algoritmi utilizzati da ATLAS per identificare i b -jet sono soggetti a continui studi volti a migliorarne le prestazioni. Un'importante modifica relativa all'algoritmo MV2 riguarda l'utilizzo di un campione di dati simulati provenienti non solo da eventi $t\bar{t}$ ma anche da $Z' \rightarrow jj$; ciò è stato pensato in modo tale da migliorare l'efficienza di b -tagging ad alti valori di p_T dei jet [37]. Tra gli eventi $t\bar{t}$ si selezionano solo quelli in cui almeno una delle due W prodotte decade leptonicamente, analogamente a quanto visto per le versioni degli algoritmi descritte nella sezione 2.5; tra gli eventi Z' si individuano solo quelli nei quali si riesce a ricostruire almeno

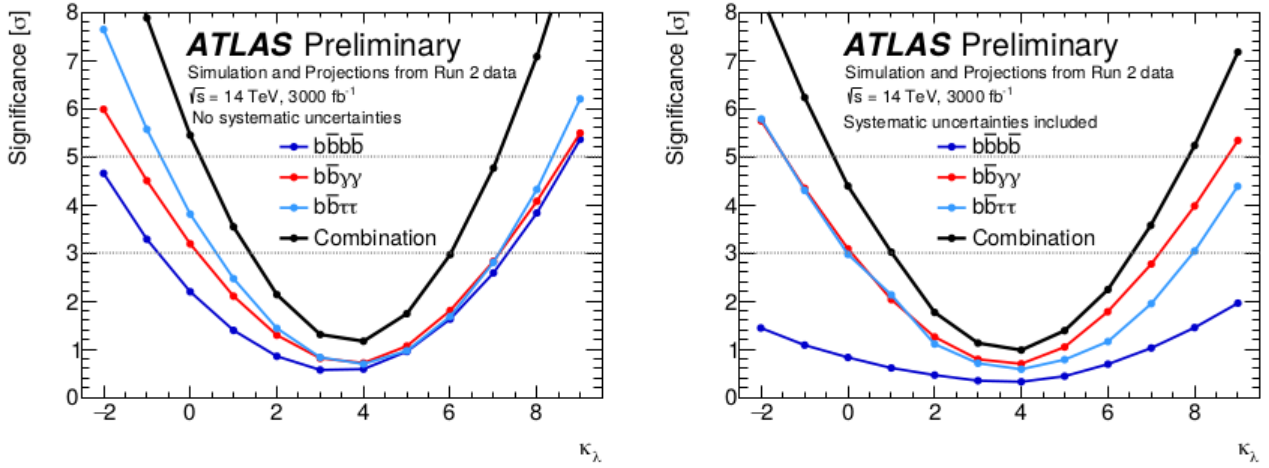


Figura 2.24: Significanza dell'eventuale osservazione di coppie di bosoni di Higgs ad HL–LHC in funzione di κ_λ nel caso in cui si considerino le sole incertezze statistiche (sinistra) o anche quelle sistematiche (destra). Da [36].

un jet. L'output di MV2 al termine delle procedure di ottimizzazione degli algoritmi di primo e secondo livello che hanno fornito i risultati presentati nella sezione 2.5 è illustrato a sinistra in figura 2.25; i valori sull'asse y sono normalizzati al numero totale di eventi.

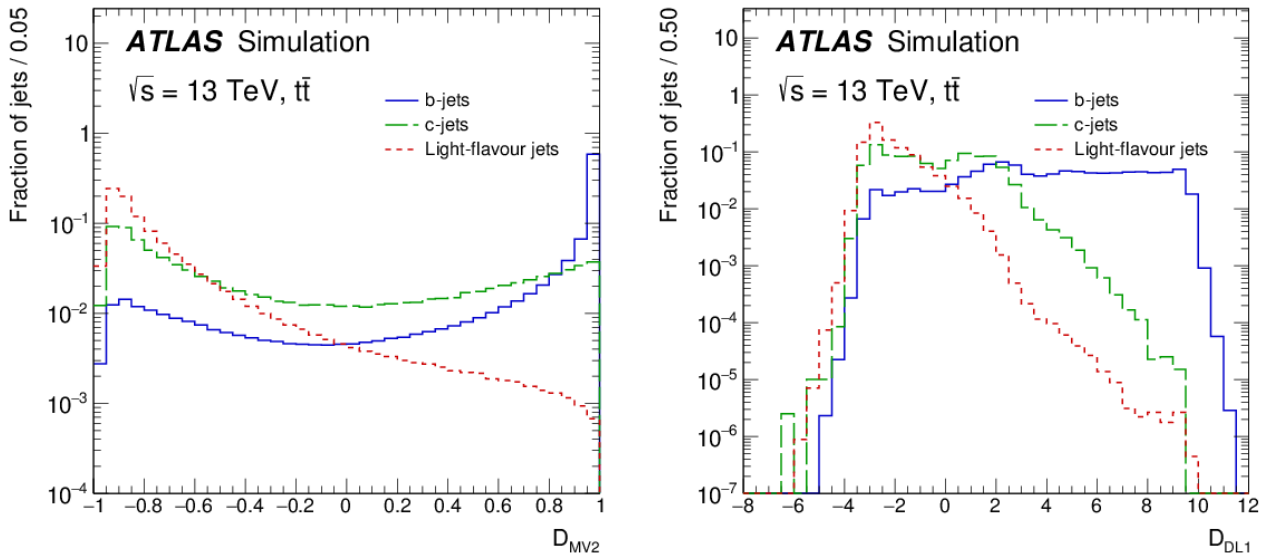


Figura 2.25: Output degli algoritmi di secondo livello MV2 (sinistra) e DL1 (destra). Da [37].

Agli algoritmi descritti nella sezione 2.5 si è affiancato nel 2017 un nuovo algoritmo di secondo livello, detto *DL1* [25]. DL1 utilizza una rete neurale basata sulle stesse variabili utilizzate come input da MV2 alle quali se ne aggiungono alcune altre ricavate dall'algoritmo JetFitter, maggiormente sensibili alla distinzione di *b*-jet da *c*-jet. Il discriminante utilizzato da DL1 è definito come

$$\ln \left(\frac{p_b}{f_c \cdot p_c + (1 - f_c) \cdot p_{\text{light}}} \right)$$

in cui p_b , p_c e p_{light} rappresentano la probabilità che il jet sia taggato come b , c o light-flavour rispettivamente e f_c indica la frazione di c -jet inclusa nel fondo. A differenza di ciò che accade nell'algoritmo MV2, il valore di f_c può essere scelto a posteriori in modo da ottimizzare le prestazioni dell'algoritmo. A destra in figura 2.25 è rappresentato l'output di DL1.

In figura 2.26 è possibile osservare un confronto tra la reiezione di light-flavour jet (sinistra) e di c -jet (destra) in funzione dell'efficienza di b -tagging ottenuta utilizzando i singoli algoritmi. Nei pannelli inferiori è rappresentato il rapporto tra le varie curve e quella relativa a MV2. Tali

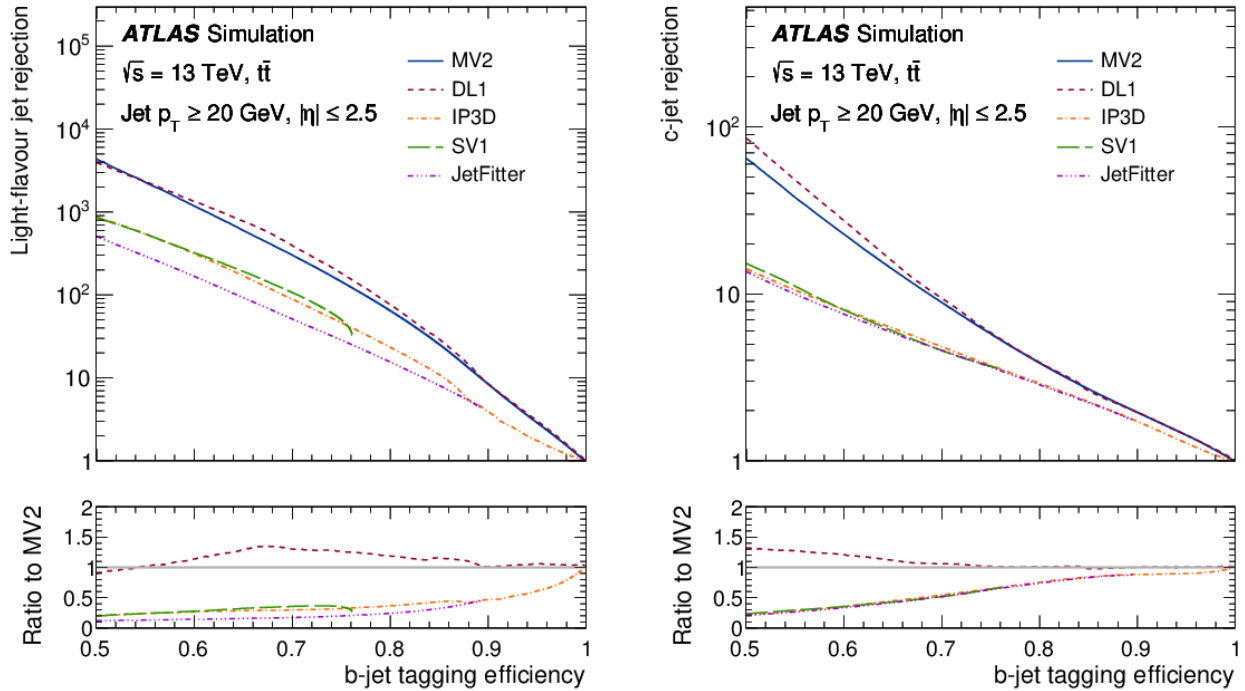


Figura 2.26: Reiezione di light-flavour jet (sinistra) e di c -jet (destra) in funzione dell'efficienza di b -tagging ottenuta dai vari algoritmi. Nei pannelli inferiori è rappresentato il rapporto tra ciascuna curva e quella relativa a MV2. Da [37].

distribuzioni dimostrano in maniera evidente l'importanza di combinare le informazioni ottenute grazie agli algoritmi di primo livello. Inoltre si nota come le prestazioni di DL1 siano migliori rispetto a quelle di MV2: considerando come punto di lavoro quello che consente di ottenere un'efficienza del 70% il miglioramento introdotto da DL1 nella reiezione di light-flavour jet e c -jet è del 30% e del 10% rispettivamente.

La calibrazione degli algoritmi di b -tagging è effettuata sulla base di un campione di eventi $t\bar{t}$ in canale dileptonico, ottenuto richiedendo l'osservazione nello stato finale di un elettrone e di un muone isolati, ovvero non circondati da altre tracce, e di due jet adronici. Tali dati sono ricavati da quelli raccolti tra il 2015 ed il 2017 ad una energia $\sqrt{s} = 13$ TeV, e corrispondono ad una luminosità integrata di 80.5 fb^{-1} . In figura 2.27 a sinistra è rappresentato un confronto tra l'efficienza di b -tagging relativa al punto di lavoro del 70% ottenuta applicando l'algoritmo MV2 alla simulazione Monte Carlo (in rosso) e ai dati (in nero) in funzione del p_T del jet. All'efficienza relativa ai dati è associato l'errore ottenuto considerando solo l'incertezza statistica (indicato con un tratto nero) oppure quella totale (indicato con una banda verde). Le principali fonti di errore non statistico provengono dall'incertezza sul modello adoperato per descrivere il processo

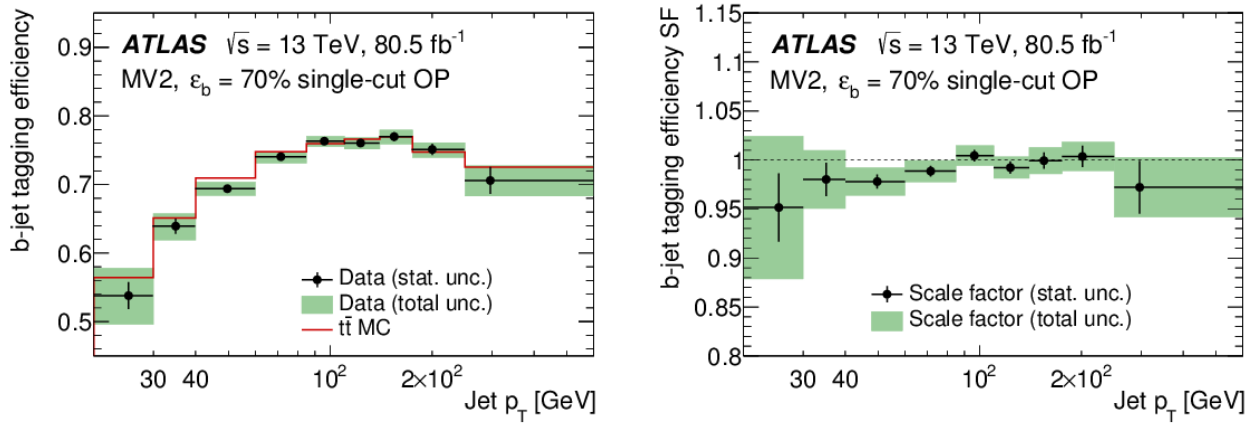


Figura 2.27: Sinistra: efficienza di b -tagging in funzione del p_T del jet ottenuta applicando l'algoritmo MV2 alla simulazione Monte Carlo (in rosso) e ai dati (in nero). Destra: fattori di scala ottenuti a partire dalle distribuzioni rappresentate a sinistra. Da [37].

di adronizzazione e dall'errore sistematico sulla scala di energia dei jet (detta *Jet Energy Scale*). La seconda predomina a bassi valori di p_T , caratterizzati dunque da errori maggiori. L'aumento dell'errore ad alti valori di momento trasverso è dovuto invece alla minore statistica disponibile. È possibile osservare come le efficienze in figura siano in accordo tra loro; ciò è confermato dai valori assunti dai fattori di scala, definiti in sezione 2.4 e rappresentati a destra in figura 2.27, globalmente compatibili con 1.

Considerazioni analoghe a quelle espone per la figura 2.27 valgono nel caso in cui l'efficienza di b -tagging ed i fattori di scala siano calcolati in funzione dell' $|\eta|$ del jet, come rappresentato in figura 2.28.

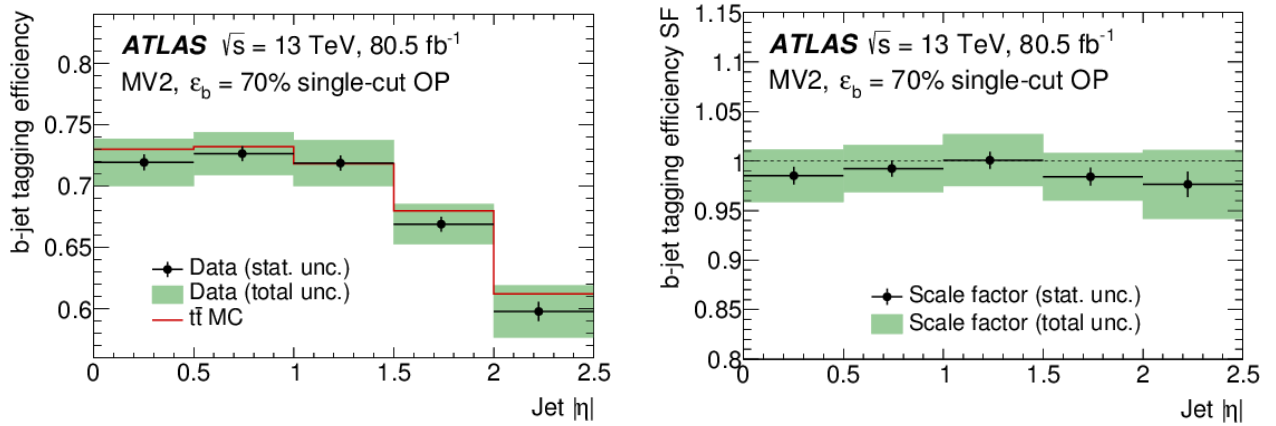


Figura 2.28: Sinistra: efficienza di b -tagging in funzione dell' $|\eta|$ del jet ottenuta applicando l'algoritmo MV2 alla simulazione Monte Carlo (in rosso) e ai dati (in nero). Destra: fattori di scala ottenuti a partire dalle distribuzioni rappresentate a sinistra. Da [37].

Capitolo 3

Rivelatori a semiconduttore

Il primo passo nell'identificazione dei b -jet è rappresentato dalla ricostruzione delle traiettorie dalle particelle nel rivelatore. L'elevato numero di particelle che si originano dal vertice di interazione protone - protone ad LHC rende necessario l'utilizzo di rivelatori di altissima risoluzione per riuscire a risolvere le numerose tracce. Le migliori prestazioni sono quelle ottenute grazie ai rivelatori a pixel di Silicio, il materiale attualmente più utilizzato per costruire dispositivi a stato solido sensibili alla radiazione ionizzante. La fisica del Silicio ha infatti conosciuto negli ultimi decenni uno sviluppo esponenziale, reso possibile anche dai progressi compiuti nell'ambito della microelettronica, che hanno permesso di ridurre le dimensioni dei pixel aumentando dunque la precisione nella ricostruzione delle tracce; inoltre, il Silicio è largamente adoperato anche in ambiti diversi da quello della fisica delle particelle.

In questo capitolo sono introdotte le proprietà generali dei rivelatori al Silicio, le modalità di formazione del segnale al passaggio di una particella ionizzante al suo interno e i danni a cui sono soggetti per effetto dell'esposizione ad elevate dosi di radiazione, come quelle presenti negli strati più interni del rivelatore di tracciamento di ATLAS. Successivamente sono descritte le proprietà principali dei rivelatori a pixel ed è illustrato un confronto tra le caratteristiche di quelli utilizzati nell'attuale Inner Detector e di quelli previsti per ITk.

3.1 Proprietà generali dei solidi

Il principio alla base della rivelazione delle particelle consiste nel raccogliere, amplificare e processare il segnale da esse generato come conseguenza dell'interazione con un mezzo. In base alle proprietà del materiale considerato tale interazione può causare diversi effetti, tra i quali la ionizzazione, sulla quale si basano ad esempio i rivelatori a gas, o l'eccitazione, sfruttata invece nel caso dei solidi.

I livelli energetici dei solidi cristallini sono organizzati in una struttura a bande. Tra queste si individuano la *banda di valenza*, alla quale appartengono gli elettroni vincolati agli atomi e pertanto non disponibili alla conduzione elettrica, e la *banda di conduzione*, ovvero quella di minore energia tra le bande che contengono elettroni non legati agli atomi, disponibili perciò alla conduzione. L'eventuale differenza energetica tra le due è indicata come *energy gap* (E_g) e definisce la *banda proibita*.

L'eccitazione del mezzo causata dalla sua interazione con una particella può stimolare la transizione di un elettrone dalla banda di valenza a quella di conduzione. Il vuoto creato nella banda di valenza da un elettrone che ha compiuto una transizione è detto *lacuna* ed ha un

comportamento assimilabile a quello di una particella carica positivamente. L'elettrone e la lacuna così generati possono ricombinarsi con cariche libere di segno opposto presenti all'interno delle rispettive bande, senza pertanto generare alcun segnale che riveli il passaggio della particella.

In presenza di un campo elettrico, invece, i portatori di carica prodotti si muovono in direzione opposta tra loro. Lacune ed elettroni sono soggetti ad un moto di deriva con diverse velocità:

$$\begin{aligned}\vec{v}_e &= -\mu_e \vec{E} \\ \vec{v}_l &= \mu_l \vec{E}\end{aligned}$$

in cui $\vec{v}_l$ indica la velocità delle lacune, $\vec{v}_e$ quella degli elettroni ed $\vec{E}$ il campo elettrico. I coefficienti μ_l e μ_e , dipendenti dalla temperatura, prendono il nome di *mobilità* e sono caratteristici del mezzo considerato. Il moto di deriva dei portatori di carica produce una corrente nel mezzo, la quale può essere misurata e processata con un'adeguata elettronica di lettura.

3.1.1 Conduttori, semiconduttori ed isolanti

La separazione energetica tra la banda di valenza e quella di conduzione varia in materiali conduttori, semiconduttori ed isolanti, come rappresentato in figura 3.1. Nei conduttori le due

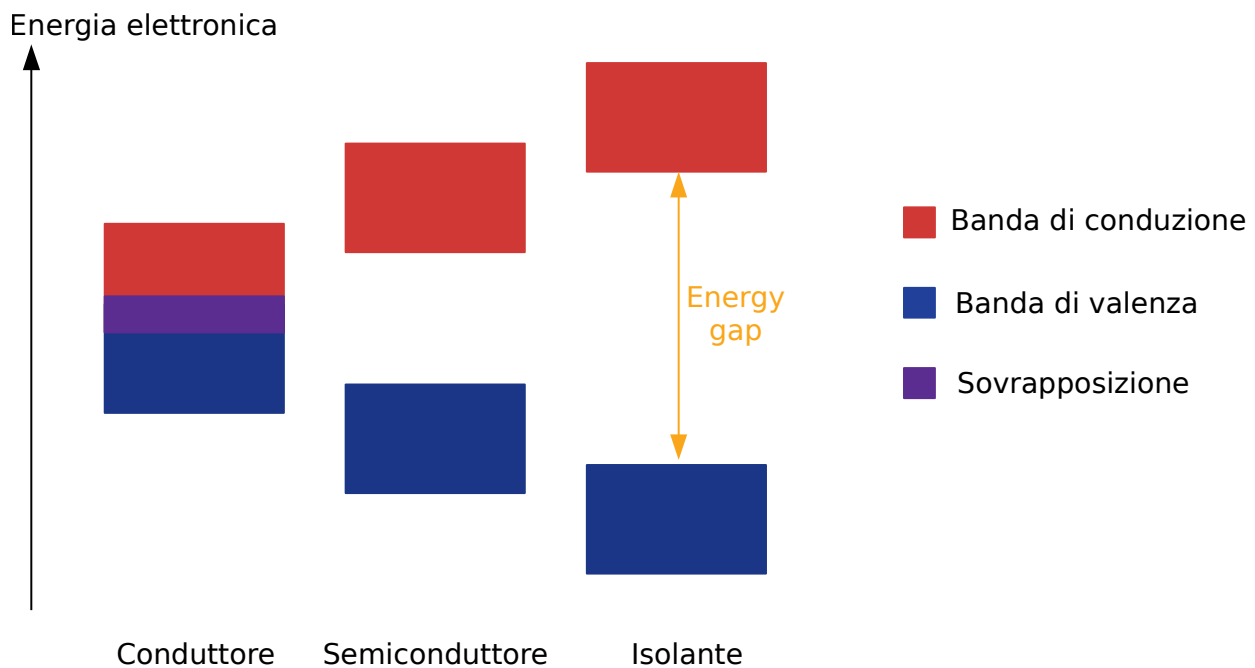


Figura 3.1: Schematizzazione della struttura a bande di conduttori, semiconduttori ed isolanti. La separazione energetica tra le due bande, se presente, è detta *energy gap*.

bande sono sovrapposte e non esiste dunque una banda proibita; la sola agitazione termica crea un grande numero di coppie elettrone - lacuna. Queste produrrebbero un elevato livello di rumore nel caso in cui si volesse rivelare il passaggio di una particella; per tale motivo i conduttori non possono essere utilizzati come rivelatori. La sovrapposizione delle bande non si verifica invece in

semiconduttori ed isolanti. Gli isolanti hanno una energy gap maggiore rispetto ai semiconduttori e di conseguenza sono caratterizzati, a parità di segnale, da una quantità minore di rumore. Tra di essi, quello maggiormente adoperato come mezzo attivo di un rivelatore di particelle è il Diamante, con $E_g = 5.5$ eV. Il suo utilizzo, tuttavia, è fortemente limitato soprattutto dalla difficoltà di sintesi del cristallo.

La quasi totalità dei rivelatori di tracciamento è realizzata con Silicio, un semiconduttore con $E_g = 1.12$ eV a temperatura ambiente ($T = 300$ K). Tale scelta è dettata dal fatto che le tecniche di lavorazione di questo materiale hanno subito notevoli progressi nel tempo, permettendo la costruzione di rivelatori ad alte prestazioni.

3.2 Rivelatori al Silicio

L'energia media necessaria per creare una coppia elettrone - lacuna nel Silicio è $E_c = 3.6$ eV; questa quantità è maggiore dell'energy gap per via del fatto che una frazione dell'energia ceduta al mezzo stimola effetti diversi dall'eccitazione (tra questi, quelli derivanti dall'interazione con gli atomi sono discussi nel paragrafo 3.4). Il valor medio della quantità di energia rilasciata da una particella ionizzante per unità di lunghezza nelle condizioni di minima ionizzazione è pari nel Silicio a $390 \text{ eV} \cdot \mu\text{m}^{-1}$. Supponendo che tale particella attraversi un rivelatore di spessore $h = 300 \mu\text{m}$ è possibile calcolare il numero medio⁽¹⁾ di coppie create:

$$\left\langle \frac{dE}{dx} \right\rangle \frac{h}{E_c} = 390 \text{ eV} \cdot \mu\text{m}^{-1} \cdot \frac{300 \mu\text{m}}{3.6 \text{ eV}} \sim 3.25 \times 10^4 \text{ coppie.} \quad (3.1)$$

Per valutare l'intensità del segnale indotto dalla deriva dei portatori di carica in un campo elettrico rispetto al rumore è necessario confrontare il numero di coppie generate come conseguenza dell'eccitazione del mezzo con quello dovuto all'agitazione termica. Tale effetto, infatti, non è trascurabile a causa del basso valore della energy gap del Silicio e determina il continuo susseguirsi di processi di generazione e ricombinazione di coppie. Il movimento dei portatori di carica presenti nel mezzo per effetto dell'agitazione termica produce una corrente, detta di *leakage*, che si sovrappone al segnale generato dalla particella. Poiché la creazione di una lacuna nella banda di valenza è dovuta alla transizione di un elettrone nella banda di conduzione i due portatori di carica hanno, in un cristallo perfetto, uguale concentrazione. Tale quantità, detta *concentrazione intrinseca* ed indicata con n_i , è fortemente dipendente dalla temperatura T :

$$n_i \propto T^{3/2} \exp\left(-\frac{E_g}{2kT}\right) \text{ cm}^{-3}$$

in cui k indica la costante di Boltzmann. Per il Silicio il valore di n_i a temperatura ambiente (alla quale $kT = 25$ meV) è $n_i \sim 1.45 \times 10^{10} \text{ cm}^{-3}$. Supponendo che il rivelatore in esame abbia una superficie $S = 10 \mu\text{m} \times 10 \mu\text{m}$ è possibile calcolare il numero medio di coppie elettrone - lacuna presenti per effetto dell'agitazione termica:

$$n_i \cdot d \cdot A = 1.45 \times 10^{10} \text{ cm}^{-3} \cdot 300 \mu\text{m} \cdot 1 \text{ cm}^2 \sim 4.35 \times 10^4 \text{ coppie.}$$

⁽¹⁾Tale valore è in realtà soggetto a delle fluttuazioni derivanti dalla natura statistica del processo di creazione di coppie elettrone - lacuna.

Tale valore è dello stesso ordine di grandezza del numero di coppie generate da una particella al minimo di ionizzazione per effetto dell'eccitazione del mezzo, indicato nell'equazione (3.1). Pertanto, per poter adoperare il Silicio come rivelatore è necessario rimuovere i portatori di carica presenti come conseguenza dell'agitazione termica; ciò è possibile realizzando una giunzione *pn*.

3.2.1 Giunzione *pn*

La sostituzione di un certo numero di atomi di un reticolo cristallino con atomi di un altro elemento, i quali prendono il nome di *impurità*, è indicata come *drogaggio* del mezzo. Si definisce *giunzione pn* l'interfaccia che si instaura all'interno di un semiconduttore nella zona di contatto tra due regioni drogate con atomi che abbiano un diverso numero di elettroni di valenza.

Il Silicio è un elemento tetravalente; pertanto, per realizzare una giunzione *pn* è necessario utilizzare atomi pentavalenti e trivalenti. Nel primo caso le impurità utilizzano quattro dei loro cinque elettroni di valenza per formare legami covalenti con il Silicio; il quinto elettrone è debolmente legato all'atomo e determina perciò un aumento del numero di elettroni liberi nel semiconduttore. Il drogaggio realizzato con atomi pentavalenti, o *donori*, è detto di tipo *n*. Viceversa, nel caso di impurità trivalenti, o *accettrici*, l'assenza del quarto legame covalente con il Silicio è interpretabile come la presenza di una lacuna; pertanto, in questo caso aumenta la concentrazione di lacune disponibili alla conduzione e si parla di drogaggio di tipo *p*.

La diversa concentrazione di elettroni e lacune che determina la formazione della giunzione origina una corrente di diffusione costituita da elettroni che si muovono dalla zona *n* a quella *p* e lacune che effettuano il movimento opposto. Una volta superata la giunzione, tali cariche si ricombinano con quelle originariamente presenti determinando una regione priva di cariche libere che prende il nome di *regione di svuotamento*. Un esempio di giunzione *pn* è rappresentato in figura 3.2.

Le distribuzioni di carica intorno alla giunzione, rappresentate da impurità ionizzate vincolate al reticolo e dunque impossibilitate al movimento, originano un campo elettrico ed un gradino di potenziale V_0 :

$$V_0 = \frac{kT}{e} \ln \left(\frac{N_D N_A}{n_i^2} \right) \quad (3.2)$$

in cui N_D indica la concentrazione degli atomi donori e N_A quella degli atomi accettori. Si può dimostrare che alla regione di svuotamento può essere applicata la formula che descrive la capacità C di un condensatore piano a facce parallele:

$$C = \frac{\epsilon_0 \epsilon_{Si} A}{W} \quad (3.3)$$

dove ϵ_0 indica la costante dielettrica del vuoto, ϵ_{Si} quella del Silicio, A la superficie della giunzione e W l'ampiezza della regione di svuotamento. Quest'ultima è determinata dalla somma tra W_n e W_p , che rappresentano l'estensione della regione di svuotamento nelle zone con drogaggio *n* ed *p* rispettivamente. Le costanti ϵ_0 ed ϵ_{Si} sono pari a

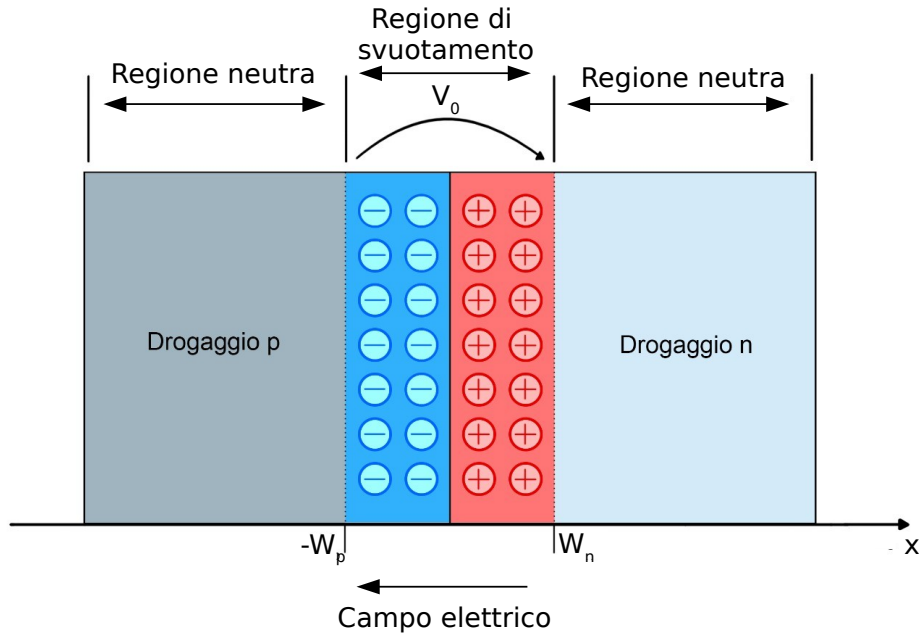


Figura 3.2: Esempio di giunzione pn. Le distribuzioni di carica intorno alla giunzione, una complessivamente positiva ed una complessivamente negativa, generano un campo elettrico ed una differenza di potenziale V_0 .

$$\epsilon_0 = 8.85 \times 10^{-12} \frac{C^2}{N \times m^2} \quad (3.4)$$

$$\epsilon_{Si} = 12.$$

Il valore di W dipende da quello di V_0 . L'espressione di W può essere ricavata risolvendo l'equazione di Poisson unidimensionale:

$$\frac{d^2 V_0}{dx^2} = -\frac{\rho(x)}{\epsilon_0 \epsilon_{Si}} \quad (3.5)$$

in cui x indica la direzione di estensione della regione di svuotamento e $\rho(x)$ la densità di carica. Quest'ultima è legata alla concentrazione delle impurità nelle due zone limitrofe alla giunzione. Ipotizzando che essa sia uniforme all'interno delle due zone e che la regione di svuotamento sia completamente priva di cariche libere, la densità di carica assume valore:

$$\rho(x) = \begin{cases} -eN_A & \text{per } -W_p < x < 0 \\ eN_D & \text{per } 0 < x < W_n. \end{cases}$$

Risolvendo la (3.5) si ottiene

$$W = W_p + W_n = \sqrt{\frac{2\epsilon_0 \epsilon_{Si}}{e} \left(\frac{1}{N_A} + \frac{1}{N_D} \right) \cdot |V_0|} \quad (3.6)$$

con

$$\begin{aligned}
W_p &= W \frac{N_D}{N_A + N_D} \\
W_n &= W \frac{N_A}{N_A + N_D}.
\end{aligned}
\tag{3.7}$$

Si osserva dunque come la larghezza della regione di svuotamento da un lato della giunzione sia legata alla concentrazione di impurità nel lato opposto.

Il valore assunto da V_0 a temperatura ambiente può essere ricavato dalla (3.2). Ipotizzando una uguale concentrazione di atomi donori e accettori, pari tipicamente a $N_A = N_D \sim 10^{15} \text{ cm}^{-3}$, si ottiene $V_0 \sim 0.5 \text{ V}$. Inserendo tale valore nella (3.6) e considerando quelli indicati nella (3.4) si ottiene $W \sim 1.63 \text{ } \mu\text{m}$. Supponendo che tale regione si estenda lungo la direzione di propagazione di una particella che attraversi lo spessore di $300 \text{ } \mu\text{m}$ precedentemente considerato risulta evidente come essa non sia ancora sufficiente per rivelare il segnale in modo efficiente. Pertanto, è necessario incrementare l'ampiezza della regione di svuotamento.

L'ampiezza della regione di svuotamento può essere modificata metallizzando le superfici del semiconduttore parallele alla giunzione ed applicandovi una tensione di alimentazione, detta di *bias* ed indicata con V_{bias} . In questo modo il gradino di potenziale assume il nuovo valore $V'_0 = V_0 + V_{bias}$, che sostituisce V_0 nella (3.6):

$$W = \sqrt{\frac{2\epsilon_0\epsilon_{Si}}{e} \left(\frac{1}{N_A} + \frac{1}{N_D} \right) \cdot |V_{bias} + V_0|}$$

La larghezza di W può aumentare o diminuire a seconda del segno di V_{bias} : per poterla incrementare, e ridurre così il numero di portatori di carica presenti per effetto dell'agitazione termica, è necessario polarizzare inversamente la giunzione, cioè applicare una tensione concorde con V_0 . Ciò è possibile collegando la regione con drogaggio p al polo negativo del generatore e quella con drogaggio n al polo positivo. In questo modo le cariche libere si allontanano dalla giunzione determinando un incremento di W .

L'aumento di W contribuisce anche alla riduzione della fluttuazione della carica Q all'interno della regione di svuotamento, assimilabile alla carica presente sulle armature di un condensatore di capacità data dalla 3.3. Tale fluttuazione è determinata dall'agitazione termica ed è quantificata dal teorema di Johnson-Nyquist, che rappresenta un caso particolare del teorema di fluttuazione-dissipazione. Secondo il teorema di Johnson-Nyquist, in un circuito elettrico caratterizzato da una capacità C , a temperatura T la fluttuazione sulla carica ha uno scarto quadratico medio pari a

$$\langle Q^2 \rangle = kTC \tag{3.8}$$

Come mostrato dall'equazione (3.3), la capacità è inversamente proporzionale all'estensione della giunzione, pertanto un aumento di W comporta una riduzione della fluttuazione della carica, indicata come *rumore elettronico*.

Il miglior rapporto segnale-rumore è ottenuto quando l'ampiezza della regione di svuotamento uguaglia lo spessore del semiconduttore. Tale condizione è ottenuta nel caso $|V_{bias}| \gg |V_0|$; un tipico valore per $W = 256 \text{ } \mu\text{m}$ è pari a $|V_{bias}| = 50 \text{ V}$ [38]. Idealmente, le uniche coppie

elettrone - lacuna presenti in questo caso sono quelle generate come conseguenza dell'eccitazione del mezzo al passaggio di una particella ionizzante; queste, per effetto del campo elettrico applicato, si muovono in direzione opposta verso i due elettrodi, su cui è indotto un segnale.

In realtà l'agitazione termica determina comunque la presenza di alcune coppie che generano una esigua corrente di leakage, di densità per unità di superficie pari a

$$J_{leak} \sim -e \frac{n_i}{\tau_g} W \sim -e \frac{n_i}{\tau_g} \sqrt{\frac{2\epsilon_0 \epsilon_{Si}}{e} \left(\frac{1}{N_A} + \frac{1}{N_D} \right) \cdot V_{bias}} \quad (3.9)$$

in cui τ_g indica il tempo medio di generazione dei portatori di carica e si è utilizzata l'approssimazione $|V_{bias} + V_0| \sim |V_{bias}|$. Tale densità ha una forte dipendenza dalla temperatura:

$$J_{leak} \sim T^2 e^{-E_g(T)/2kT}$$

Per questo motivo nei rivelatori basati sul Silicio può essere necessario introdurre un sistema di raffreddamento che consenta di stabilizzare la densità della corrente di leakage. Un valore tipico di corrente di leakage è pari a ~ 100 pA per un rivelatore di dimensioni $50 \mu\text{m} \times 400 \mu\text{m} \times 256 \mu\text{m}$ in corrispondenza di una temperatura di 0°C [38].

La relazione tra corrente e tensione attraverso una giunzione pn , spesso indicata come *curva I-V*, è rappresentata in figura 3.3; la configurazione di polarizzazione inversa è relativa al caso $V_{bias} < 0$. È evidente come all'applicazione della tensione di polarizzazione inversa si generi

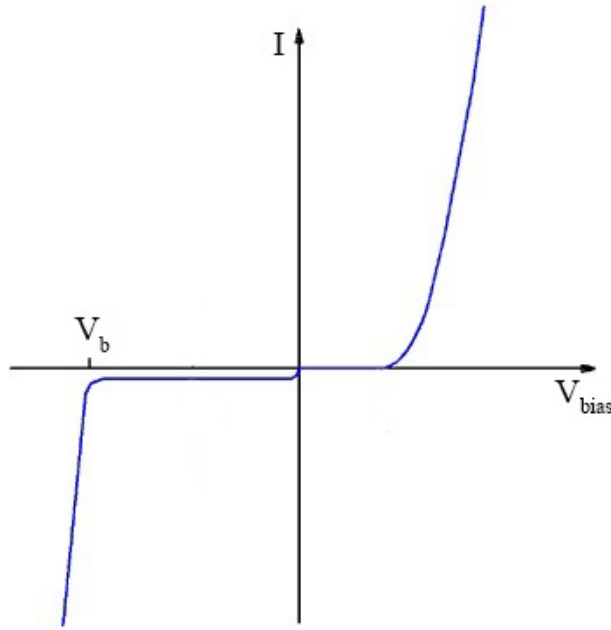


Figura 3.3: Relazione tra tensione applicata ai capi di una giunzione pn e corrente che fluisce al suo interno.

la corrente di leakage, proporzionale a $\sqrt{|V_{bias}|}$; una volta raggiunta la tensione che permette di estendere la regione di svuotamento all'intero spessore del rivelatore la corrente assume un andamento pressoché costante. Tale comportamento della corrente di leakage è espresso dalla relazione (3.9). In corrispondenza del valore V_b si verifica il *breakdown* del dispositivo, caratterizzato da un aumento repentino della corrente. Questo comportamento è dovuto alla

concomitanza dell'effetto Zener e dell'effetto valanga. Il primo indica la generazione di un elevato flusso di elettroni verso la banda di conduzione per effetto tunnel, favorito dall'aumento del campo elettrico; nel secondo caso gli elettroni liberi hanno acquisito una energia tale da generare ulteriori cariche disponibili per la conduzione per effetto di urti con gli atomi del reticolo. Un tipico valore di V_b è pari a circa -350 V per sensori planari e -50 V per sensori 3D [30]; la loro defizione è fornita nella sezione 3.5.2.

La regione $V_{bias} > 0$ è relativa al caso di polarizzazione diretta, in cui la tensione è applicata in direzione opposta rispetto al caso di polarizzazione inversa. In questa circostanza l'ampiezza della regione di svuotamento diminuisce; tale modalità, pertanto, non è di interesse per la realizzazione di rivelatori al Silicio.

3.3 Formazione del segnale

Il fatto che la corrente sia generata dal moto di deriva dei portatori di carica implica che questa si formi prima che tutte le coppie prodotte dall'eccitazione del mezzo siano raccolte dagli elettrodi.

La corrente istantanea i indotta su un elettrodo per effetto della deriva di una carica q con velocità $\vec{v}$ può essere calcolata utilizzando il teorema di Shockley-Ramo [39, 40]:

$$i = q\vec{v} \cdot \vec{E}_w. \quad (3.10)$$

$\vec{E}_w$ prende il nome di *campo di weighting* e rappresenta il campo elettrico che si produrrebbe tra i due elettrodi nel caso in cui quello in esame fosse collegato al polo positivo di un generatore che fornisce una tensione unitaria e tutti gli altri⁽²⁾ a massa. La carica indotta dal moto di q dalla posizione $\vec{x}_1$ a quella $\vec{x}_2$ nell'intervallo compreso tra gli istanti t_1 e t_2 è pari all'integrale della (3.10):

$$Q = \int_{t_1}^{t_2} i(t) dt = q[\phi_w(\vec{x}_1) - \phi_w(\vec{x}_2)]$$

in cui ϕ_w è detto *potenziale di weighting*, associato ad $\vec{E}_w$ ed ottenuto risolvendo l'equazione di Poisson.

Nel caso in cui la larghezza dell'elettrodo di raccolta della carica sia paragonabile allo spessore della regione attiva il potenziale ϕ_w è una funzione lineare della distanza dall'elettrodo. Ciò comporta che la carica istantanea indotta su di esso dipenda solo dalla posizione di q e non dal valore della tensione applicata o dalla distribuzione geometrica dei portatori di carica.

La linearità del potenziale di weighting, tuttavia, non è più valida se l'elettrodo di raccolta della carica ha una superficie inferiore rispetto alla sezione della regione attiva (come illustrato nella sezione 3.5, questo vale nei rivelatori a pixel). In tal caso la maggior parte del segnale è indotto in prossimità dell'elettrodo, nella fase finale del moto di deriva delle cariche; il contributo di quelle che si allontanano da esso diventa trascurabile.

⁽²⁾Il teorema di Ramo è valido nel caso generale di N elettrodi.

3.4 Danni da radiazione

Durante la loro propagazione all'interno del semiconduttore le particelle possono interagire non solo con gli elettroni atomici ma anche con gli atomi del reticolo. Ciò può determinare il loro spostamento dalla posizione originaria e dunque un cambiamento delle proprietà del mezzo. Gli atomi non più vincolati al reticolo possono muoversi all'interno del materiale per occupare un'altra posizione e sono indicati come *difetti*; la loro interazione con altri atomi può causare la creazione di ulteriori difetti.

Una rottura della periodicità del reticolo determina la formazione di nuovi livelli energetici all'interno della banda proibita; questi fungono da centri di generazione e ricombinazione delle cariche, causando un aumento della corrente di leakage. Inoltre i difetti, essendo elettricamente carichi, possono portare ad un'inversione del drogaggio originario del mezzo. Infine, il vuoto determinato dalla rimozione di un atomo dal reticolo cristallino può fungere da centro di cattura della carica generata dal passaggio di una particella. Tale effetto è particolarmente grave nel caso di elettrodi di piccola dimensione, nei quali il segnale è dovuto alle cariche che riescono a raggiungere le vicinanze dell'elettrodo (come illustrato nella sezione 3.3).

La deformazione del reticolo causata dall'interazione delle particelle con gli atomi è particolarmente probabile in presenza di elevate dosi di radiazione, come quelle generate in prossimità del punto di interazione tra i fasci accelerati ad LHC. Per questo motivo gran parte delle attività di ottimizzazione dei rivelatori di tracciamento è finalizzata all'aumento della loro resistenza ai danni indotti dalla radiazione.

3.5 Rivelatori a pixel

I primi rivelatori di tracciamento basati sul Silicio sono stati quelli a strip. Questi sono ottenuti drogando selettivamente la superficie di un supporto di Silicio precedentemente drogato in modo opposto, detto *p-bulk* o *n-bulk* in base alla natura delle impurità utilizzate; un esempio di rivelatore a strip è rappresentato a sinistra in figura 3.4. Ciascuno degli impianti realizzati, detto

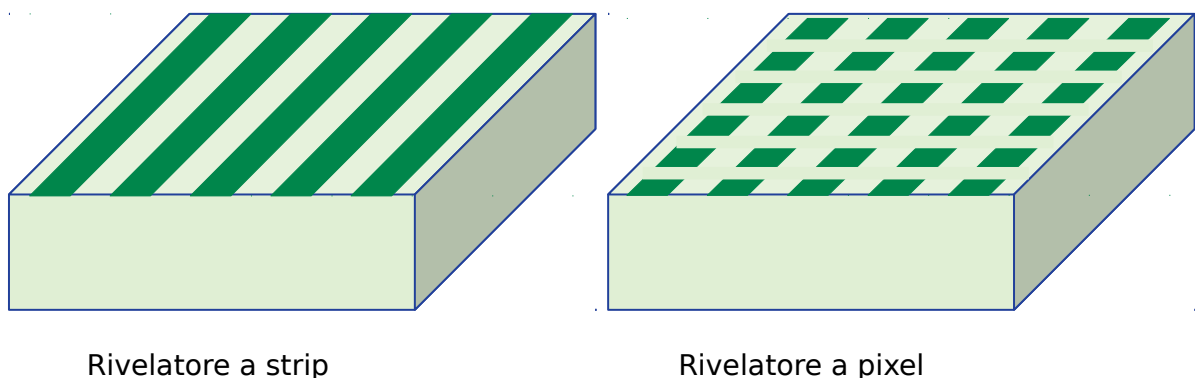


Figura 3.4: Rappresentazione schematica di un rivelatore a strip (sinistra) e di uno a pixel (destra).

strip, origina una regione di svuotamento che si estende in una delle due direzioni del piano del rivelatore; polarizzandoli inversamente in modo da estendere la regione di svuotamento all'intera

sezione del bulk si può utilizzare il dispositivo realizzato per rivelare particelle. La segmentazione a strip permette di ricostruire una delle coordinate (quella perpendicolare alla direzione degli impianti) alle quali è avvenuto il passaggio della particella nel rivelatore. Uno dei metodi adoperati per ricavare informazioni bidimensionali consiste nel realizzare delle strip anche sulla superficie inferiore del bulk, con un angolo non nullo rispetto alle prime. I rivelatori così ottenuti, tuttavia, forniscono risultati ambigui nel caso in cui siano attraversati contemporaneamente da più particelle. Si consideri, a titolo esemplificativo, il caso in cui le strip siano ortogonali tra loro, come rappresentato in figura 3.5. Le interazioni di due particelle nelle posizioni indicate in rosso

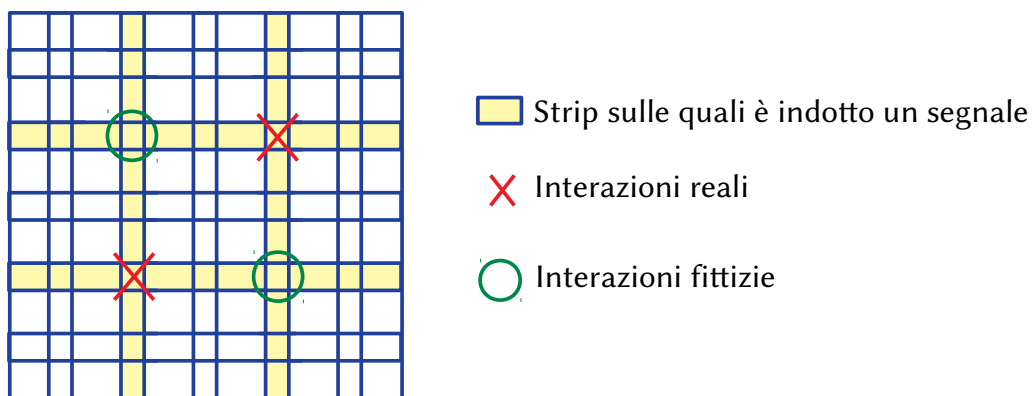


Figura 3.5: Illustrazione di un caso di ambiguità nella determinazione dei punti di interazione di due particelle con un rivelatore a strip.

determina l'induzione di un segnale sulle quattro strip evidenziate in giallo. La loro intersezione indica come possibili punti di interazione della particella non solo quelli reali ma anche quelli evidenziati con un cerchio verde.

La necessità di risolvere tale problema ha determinato lo sviluppo dei rivelatori a pixel; un esempio è rappresentato a destra in figura 3.4. Questi sono costituiti da una matrice di giunzioni *pn* polarizzate inversamente che consente di ottenere informazioni bidimensionali riguardo alle traiettorie descritte dalle particelle evitando le ambiguità rappresentate in figura 3.5. Questo è il motivo principale per cui i rivelatori a pixel sono adoperati negli strati più interni dei tracciatori, caratterizzati da un elevato flusso di particelle.

La regione attiva del rivelatore, ovvero quella in cui avviene la generazione del segnale, è indicata come *sensore*. L'insieme dei circuiti che leggono, amplificano e digitalizzano il segnale prodotto è detto *front-end*. Il tipo di collegamento tra le due componenti definisce la distinzione tra rivelatori ibridi e monolitici. In base alla geometria degli elettrodi, inoltre, si distingue tra rivelatori planari e 3D.

3.5.1 Rivelatori ibridi

I rivelatori attualmente più utilizzati sono quelli *ibridi*, in cui l'elettronica di front-end ed il sensore sono realizzati separatamente e successivamente connessi tra loro.

I circuiti di front-end sono integrati in un chip sul quale sono saldate delle sfere conduttive del diametro dell'ordine della decina di μm , dette *bump-bond*, utilizzate per il collegamento elettrico

con i pixel del sensore. In figura 3.6 è rappresentato un esploso di un rivelatore ibrido costruito utilizzando questa tecnica, nota come *bump-bonding*.

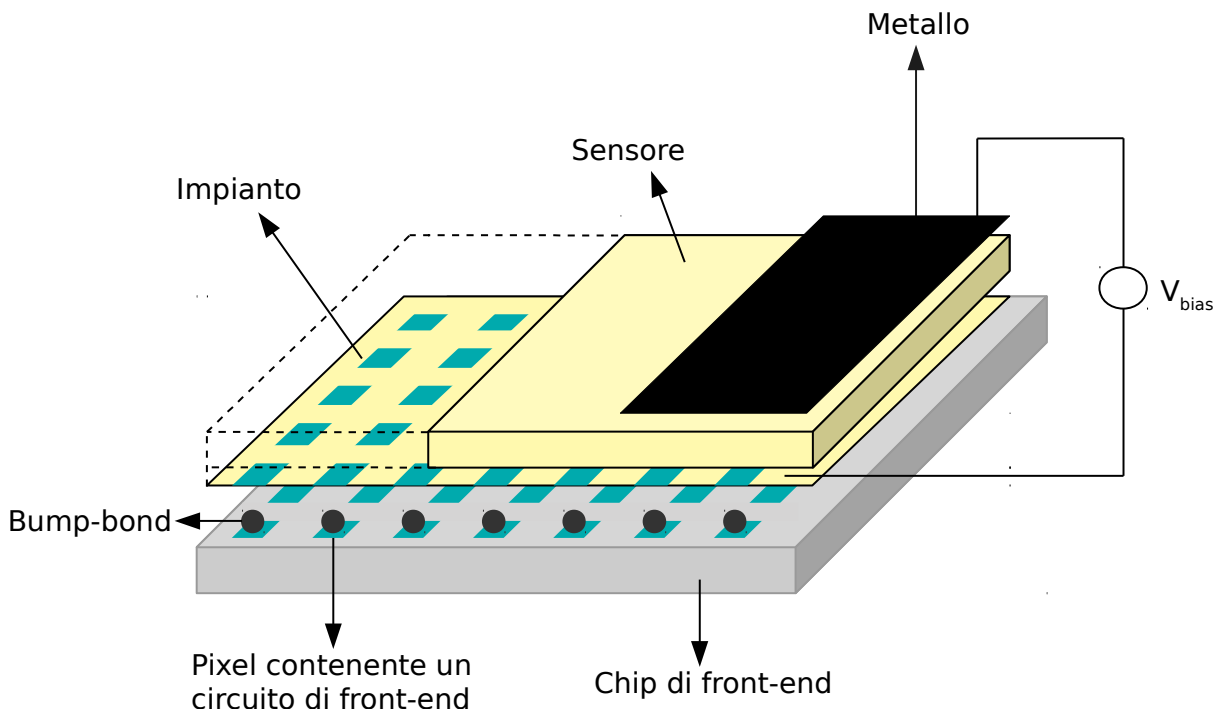


Figura 3.6: Esploso di un rivelatore ibrido realizzato tramite la tecnica del bump-bonding. Sono visibili il chip di front-end, i bump-bond ed il sensore, rivestito con uno strato di alluminio per l'applicazione della tensione di bias.

Nell'immagine raffigurata ogni pixel del sensore è collegato ad uno dei circuiti di front-end, i quali sono organizzati in una matrice di pixel uguale a quella del sensore. La posizione dei bump-bond può essere ottimizzata in modo da connettere lo stesso chip a sensori con pixel di diverse dimensioni. In figura 3.7, ad esempio, è illustrato come lo stesso chip possa essere utilizzato per la lettura di sensori con pixel di dimensioni $50\ \mu\text{m} \times 50\ \mu\text{m}$ (a sinistra), $100\ \mu\text{m} \times 25\ \mu\text{m}$ (al centro) oppure $100\ \mu\text{m} \times 100\ \mu\text{m}$ (a destra). L'area del pixel del circuito

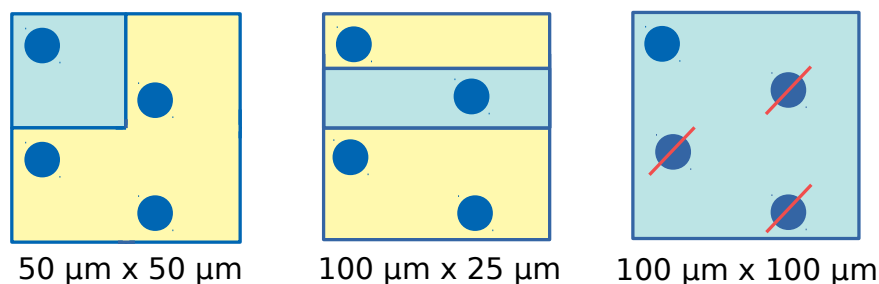


Figura 3.7: Connessione tra lo stesso chip di front-end e sensori con pixel di diverse dimensioni.

di front-end in tutti e tre i casi è la stessa; in celeste è indicata la superficie del pixel del sensore.

I cerchi blu rappresentano i bump-bond; quelli non utilizzati per il collegamento al sensore sono cancellati con una linea rossa. Questa versatilità di connessione tra pixel del chip e del sensore consente di limitare il numero di geometrie di chip che devono essere costruiti ed ottimizzati, con conseguente riduzione di tempi e costi.

Il dispositivo ottenuto connettendo un sensore all'elettronica di front-end è indicato nel seguito con il termine *modulo*; questo può essere realizzato utilizzando più chip per la lettura di un unico sensore.

La tecnologia ibrida offre la possibilità di ottimizzare separatamente il chip di front-end ed il sensore e ciò ha determinato il suo rapido sviluppo. Tuttavia, questo tipo di rivelatore presenta alcuni svantaggi. L'utilizzo di sfere metalliche per il collegamento all'elettronica di front-end determina un'aumento della capacità complessiva del sensore e dunque del rumore elettronico, come indicato dall'equazione (3.8). Inoltre, la realizzazione dell'elettronica su un chip separato dal sensore comporta un aumento del materiale inattivo con cui le particelle che si propagano nel rivelatore devono interagire. Infine, la realizzazione dei bump-bond è una procedura complicata e soggetta ad un'elevata probabilità di insuccesso in fase di produzione dei moduli.

I problemi dei rivelatori ibridi sono risolti grazie ad una tecnologia indicata come *monolitica*, attualmente in fase di sviluppo, in cui l'elettronica di front-end è integrata all'interno del sensore; tuttavia tali rivelatori sono, allo stato attuale, meno resistenti rispetto a quelli ibridi ai danni indotti dall'interazione con elevate dosi di radiazioni. Maggiori informazioni riguardo ai rivelatori monolitici possono essere trovate in Appendice C.

3.5.2 Rivelatori planari e 3D

I primi sensori a pixel ad essere stati realizzati sono quelli detti *planari*, in cui gli elettrodi di raccolta della carica sono organizzati in una matrice bidimensionale sulla superficie di un bulk, come rappresentato a destra in figura 3.4.

Sulla base del tipo di drogaggio del bulk e degli impianti i rivelatori planari possono essere distinti in varie categorie, caratterizzate da diverse proprietà. I rivelatori planari previsti per ITk saranno realizzati con impianti di tipo n^+ su un bulk di tipo p ; il drogaggio n^+ è ottenuto utilizzando una maggiore concentrazione di impurezze rispetto ad uno di tipo n . Tale geometria è indicata come *n-in-p*; in figura 3.8 è rappresentata una sezione di una cella di rivelazione, incluso il pixel di front-end nel chip di lettura. Il drogaggio n^+ permette di estendere la regione di svuotamento soprattutto nel bulk, come indicato dalle relazioni (3.7). L'impianto n^+ è rivestito da uno strato metallico sul quale è saldato un bump-bond per la connessione con il chip di front-end. All'estremità inferiore del p -bulk è realizzato un impianto di tipo p^+ , rivestito di uno strato metallico per il collegamento al generatore di tensione.

Una tipologia più recente di rivelatori a pixel, rappresentata a sinistra in figura 3.9, è quella detta *3D*. In questa geometria gli elettrodi sono costituiti da cilindri di Silicio drogati di tipo n^+ e p^+ impiantati perpendicolarmente alle superfici del bulk, il quale può essere di tipo n oppure p . Analogamente al caso dei rivelatori planari, anche quelli 3D utilizzano prevalentemente bulk di tipo p . Il numero di elettrodi di tipo n^+ circondati da quelli p^+ all'interno di ogni pixel determina una diversa denominazione del rivelatore, detto 1E nel caso in cui ogni pixel contenga un unico elettrodo, 2E nel caso in cui ne contenga due e così via. Una rappresentazione dall'alto di un pixel di un sensore 2E è rappresentata a destra in figura 3.9.

La geometria 3D è stata sviluppata più recentemente rispetto a quella planare, ma ha subito trovato impiego nei rivelatori di particelle di alta energia grazie ad alcuni vantaggi derivanti dalla

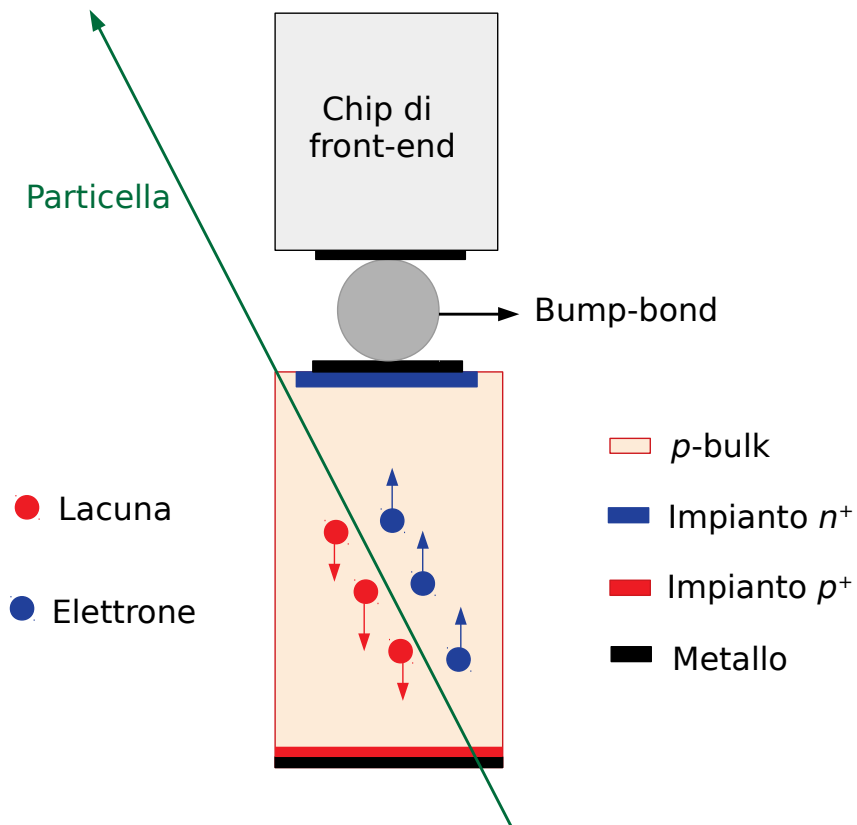


Figura 3.8: Sezione di una cella di rivelazione di un rivelatore ibrido planare e di un pixel di front-end.

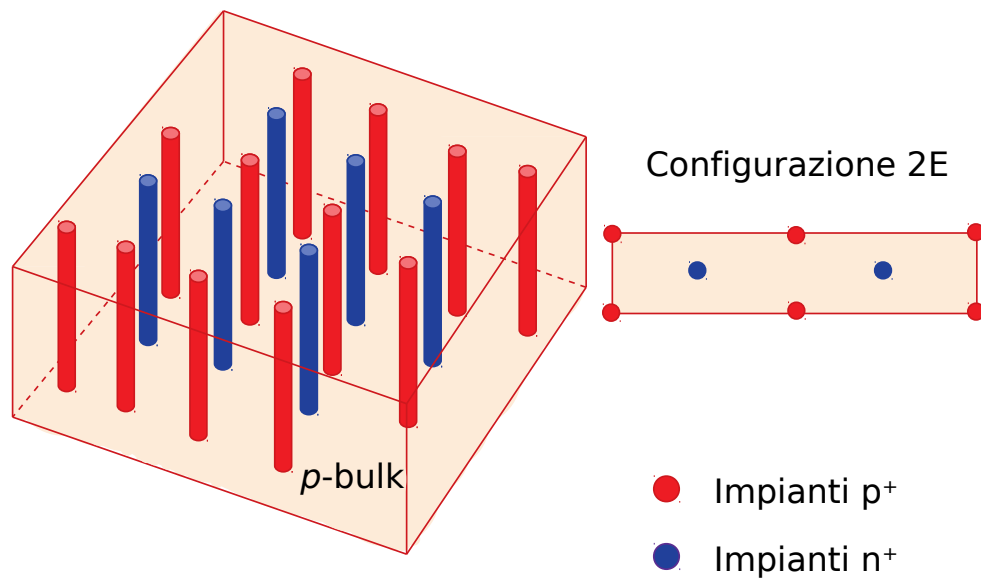


Figura 3.9: Sinistra: rappresentazione schematica di un rivelatore a pixel 3D, in cui gli elettrodi sono impiantati perpendicolarmente alla superficie del bulk. Destra: rappresentazione dall'alto di un pixel di un sensore 2E, contenente due elettrodi di tipo n^+ .

differente disposizione degli elettrodi; parte dell'IBL dell'attuale Inner Detector di ATLAS, ad esempio, è stata realizzata usando rivelatori 3D con bulk di tipo p . Il movimento dei portatori di carica avviene parallelamente alle superfici del sensore, come rappresentato in figura 3.10 e diversamente da quanto illustrato in figura 3.8 per i rivelatori planari. La distanza tra gli elet-

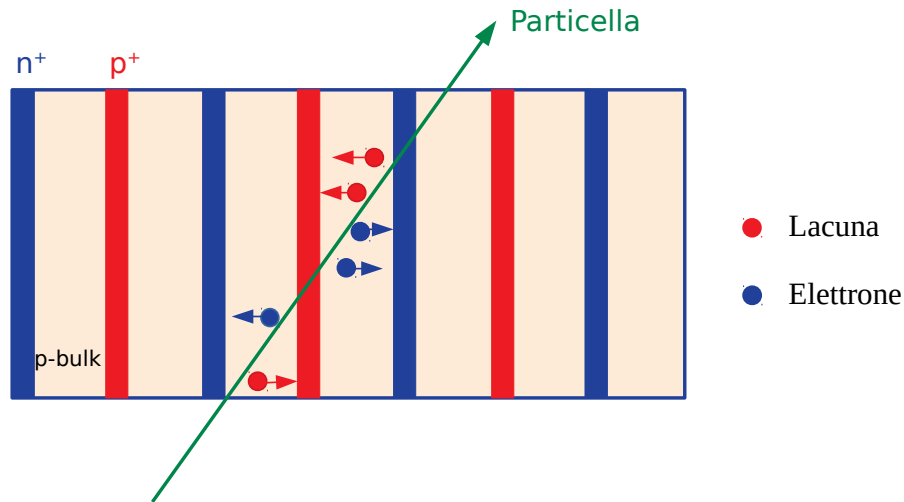


Figura 3.10: Illustrazione del movimento di deriva dei portatori di carica generati per effetto dell'eccitazione del mezzo all'interno di un rivelatore a pixel 3D.

trodi, non più vincolata dallo spessore del sensore, è minore rispetto a quella caratteristica della geometria planare; ciò determina una più rapida formazione del segnale ed una diminuzione della tensione necessaria per rimuovere le cariche libere presenti per effetto dell'agitazione termica. Inoltre, la geometria 3D permette di limitare le conseguenze del danneggiamento del rivelatore dovuto alla radiazione alla quale è sottoposto. Come spiegato nella sezione 3.4, infatti, l'interazione di particelle con gli atomi del reticolo può originare dei centri di cattura delle cariche prodotte come conseguenza dell'eccitazione del mezzo, determinando una diminuzione del segnale indotto sugli elettrodi. Dal momento che la cattura delle cariche dipende dalla distanza da queste percorsa prima di raggiungere l'elettrodo di raccolta è evidente come si verifichi in misura minore nei rivelatori 3D, nei quali c'è pertanto una maggiore correlazione tra il segnale generato e l'energia depositata da una particella all'interno del sensore.

La geometria 3D, tuttavia, presenta anche alcuni svantaggi; tra questi, le maggiori difficoltà di produzione ed una diminuzione del volume disponibile alla generazione dei portatori di carica dovuto alla presenza degli elettrodi all'interno del bulk. Per ridurre l'ultima problematica alcuni sensori sono realizzati utilizzando degli elettrodi che penetrano solo parzialmente lo spessore del bulk.

3.6 Confronto tra i rivelatori a pixel dell'Inner Detector e di ITk

Come discusso nel capitolo 2, l'Inner Detector di ATLAS sarà sostituito da un nuovo rivelatore di tracciamento, detto ITk, il quale entrerà in funzione a partire dalla metà del 2027. L'esigenza

di resistere all'elevata dose di radiazione conseguente all'aumento di luminosità previsto ad HL–LHC ha motivato il cambiamento di alcuni dei parametri caratteristici dei rivelatori a pixel utilizzati nei due casi. Nelle sezioni successive, relative rispettivamente ai rivelatori 3D e planari, è rappresentato un confronto tra tali parametri.

Una caratteristica comune all'Inner Detector e ad ITk è l'utilizzo di rivelatori ibridi; tale scelta è stata motivata nella sezione 3.5.1.

3.6.1 Sensori 3D

Le principali caratteristiche dei sensori 3D utilizzati nell'Inner Detector e di quelli previsti per ITk sono confrontate in tabella 3.1.

Sensori 3D		
Rivelatore	Inner Detector	Inner Tracker
Collocazione	Parte di IBL	Primo strato
Drogaggio bulk	p	p
Spessore bulk [μm]	230	150
Dimensione pixel [$\mu\text{m} \times \mu\text{m}$]	50×250	25×100 (barrel) - 50×50 (end-cap)
Configurazione elettrodi n^+	2E	1E

Tabella 3.1: Confronto tra i parametri caratteristici dei sensori 3D attualmente utilizzati nell'Inner Detector e di quelli previsti per ITk. Nel caso di pixel rettangolari, il lato più corto è orientato lungo la direzione di propagazione dei fasci.

I sensori 3D sono collocati nello strato più interno di entrambi i rivelatori, il quale riceve una dose di radiazione maggiore rispetto ai successivi. Ciò è motivato dalla maggiore resistenza alla radiazione dei rivelatori 3D rispetto a quelli planari discussa nella sezione 3.5.2.

Nell'Inner Tracker saranno utilizzati sensori 3D con pixel di dimensioni ridotte rispetto a quelli installati nell'Inner Detector. Poiché la risoluzione spaziale ottenuta nel tracciamento migliora al diminuire delle dimensioni del pixel, tale cambiamento permetterà di bilanciare il peggioramento dovuto alla presenza di un maggior numero di particelle da ricostruire derivante dall'aumento di luminosità; le ripercussioni previste sull'efficienza di b -tagging sono state discusse nel capitolo 2. La dimensione dei pixel 3D da utilizzare nell'ITk, tuttavia, deve essere ancora stabilita sulla base dei risultati delle simulazioni delle prestazioni del rivelatore ottenute considerando le due possibilità indicate in tabella 3.1.

Il numero di elettrodi impiantati in ogni pixel dei sensori utilizzati nell'IBL è pari a 2; questa configurazione è maggiormente resistente alla radiazione rispetto a quella 1E grazie alla minore distanza percorsa dalle cariche. Tuttavia, la configurazione 1E è di più semplice fabbricazione, come conseguenza della maggiore distanza tra gli elettrodi; per questo motivo è stata selezionata per la realizzazione dei sensori da utilizzare nell'Inner Tracker.

3.6.2 Sensori planari

Le principali caratteristiche dei sensori planari utilizzati nell'Inner Detector e di quelli previsti per ITk sono confrontate in tabella 3.2. In quanto meno resistenti alla radiazione, sono utilizzati sia nell'Inner Detector che in ITk in strati successivi al primo.

Sensori planari			
Rivelatore	Inner Detector		Inner Tracker
Collocazione	Parte di IBL	Strati successivi	Strati successivi
Drogaggio	n^+ -in- n	n^+ -in- n	n -in- p
Spessore bulk [μm]	150	256	100 e 150
Dimensione pixel [$\mu\text{m} \times \mu\text{m}$]	50×250	50×400	50×50

Tabella 3.2: Confronto tra i parametri caratteristici dei sensori planari attualmente utilizzati nell'Inner Detector e di quelli previsti per ITk. Nel caso di pixel rettangolari, il lato più corto è orientato nella direzione in cui si richiede la maggiore precisione.

I sensori utilizzati nell'IBL sono costituiti da impianti di tipo n^+ su un n -bulk; in questo modo la giunzione pn si crea all'interfaccia tra il bulk e l'impianto p^+ sottostante. L'utilizzo di un bulk di tipo p , come quello scelto per ITk, permette di ottenere maggiori velocità di formazione del segnale rispetto ad un rivelatore realizzato con un n -bulk. Infatti, come discusso nella sezione 3.3, la formazione del segnale in un rivelatore a pixel avviene nella fase finale del moto di deriva delle cariche verso gli elettrodi. In un rivelatore p -in- n pertanto il segnale è generato dalle lacune che si trovano in prossimità dell'elettrodo p , collegato al polo negativo del generatore di tensione in modo da polarizzare inversamente la giunzione pn . Le lacune, tuttavia, hanno una mobilità inferiore rispetto a quella degli elettroni (in quanto inversamente proporzionale alla loro massa) e pertanto la formazione del segnale è più lenta rispetto a quella dei pixel n -in- p . Per questo motivo rivelatori n -in- p sono maggiormente adatti ad essere utilizzati ad HL-LHC, quando l'elevato numero di particelle generate dall'interazione tra i fasci di protoni richiederà un aumento della velocità di processamento del segnale.

Un'altra differenza tra i rivelatori planari utilizzati nell'Inner Detector e quelli previsti per ITk è rappresentata da una diminuzione dello spessore del materiale attivo. La riduzione della distanza percorsa dalle cariche durante il loro moto di deriva verso gli elettrodi comporta una diminuzione del segnale indotto, ma anche una serie di vantaggi. Il segnale è generato in tempi minori, rendendo il rivelatore più veloce. Inoltre, diminuisce la possibilità di cattura delle cariche per effetto della presenza di imperfezioni del reticolo; i rivelatori più sottili sono dunque maggiormente resistenti alla radiazione. Infine, la riduzione della differenza tra lo spessore del bulk e la larghezza del pixel comporta un aumento della regione di linearità del potenziale di weighting; di conseguenza, le cariche distanti dall'elettrodo contribuiscono maggiormente all'induzione del segnale su di esso. Ciò permette di bilanciare la perdita di segnale dovuta alla diminuzione della distanza percorsa dalle cariche.

Analogamente ai rivelatori 3D di ITk, anche quelli planari avranno dei pixel di dimensione inferiore rispetto a quelli dell'Inner Detector e ciò contribuirà a mantenere un'elevata risoluzione nel tracciamento anche a valori di luminosità più elevati.

Capitolo 4

Il chip RD53A

L'aumento di luminosità previsto ad HL—LHC rende necessaria la sostituzione dei rivelatori a pixel e a strip dell'Inner Detector, non adatti a sostenere tale cambiamento. Nel 2015 è stata approvata la realizzazione di un prototipo dei chip da utilizzare per i pixel dell'ITk, detto RD53A. A partire dal 2017 tale chip è stato distribuito a diversi gruppi appartenenti alla collaborazione ATLAS affinché questi potessero svolgere attività di ottimizzazione e confrontare i risultati dei loro studi.

In questo capitolo sono illustrate le caratteristiche principali del chip RD53A, sul quale sono state svolte le attività di misura e di analisi oggetto della presente tesi. Inoltre, sono discussi i risultati ottenuti dai vari gruppi di ricerca e i prossimi passi verso la realizzazione del chip finale, da utilizzare in ITk.

4.1 Caratteristiche generali

I chip da utilizzare nell'Inner Tracker dovranno essere maggiormente resistenti alla radiazione rispetto a quelli attualmente presenti nell'Inner Detector, dovranno avere una maggiore granularità per adattarsi alla diminuzione della dimensione dei pixel del sensore e dovranno essere in grado di gestire un aumento della quantità di dati da processare. I requisiti che tali chip dovranno rispettare sono dettati dallo strato più interno del rivelatore a pixel, il quale riceve il maggior flusso di particelle e quindi anche la maggiore dose di radiazioni.

Nel 2013 è nata la collaborazione RD53 [41], con l'obiettivo di progettare un chip di lettura adatto ad essere utilizzato per i rivelatori a pixel ibridi di ATLAS e CMS. Il primo prototipo di chip sviluppato dalla collaborazione è RD53A, con il quale si è voluta dimostrare la possibilità di soddisfare i requisiti di progetto con un circuito integrato di grandi dimensioni realizzato con tecnologia CMOS 65 nm. Questa rappresenta un'evoluzione delle tecnologie adoperate per la realizzazione dei chip di lettura dell'IBL (CMOS 130 nm) e degli altri strati del rivelatore a pixel dell'Inner Detector (CMOS 250 nm), caratterizzata dalla diminuzione delle dimensioni dei transistor utilizzati. Ciò si traduce nella possibilità di integrare l'elettronica di front-end in pixel di dimensioni minori pur mantenendo (o addirittura estendendo) le funzionalità analogiche e digitali dei chip predecessori, soddisfacendo dunque la richiesta di aumento di granularità.

In quanto prototipo, RD53A non è stato costruito rispettando tutte le specifiche previste per il chip finale, ed in esso sono implementate varianti di elettronica di front-end che rendono non uniforme la matrice dei pixel. Un altro obiettivo del chip, infatti, è quello di confrontare le prestazioni delle diverse logiche di front-end implementate, in modo da selezionare quella da

utilizzare per la realizzazione dei chip che saranno effettivamente utilizzati da ATLAS e CMS durante la fase di HL–LHC.

4.2 Struttura del chip

Il chip RD53A ha una larghezza di 20 mm ed un'altezza di 11.5 mm; la sua struttura è rappresentata in figura 4.1. Si compone di una matrice di pixel e due periferie poste sui due lati più

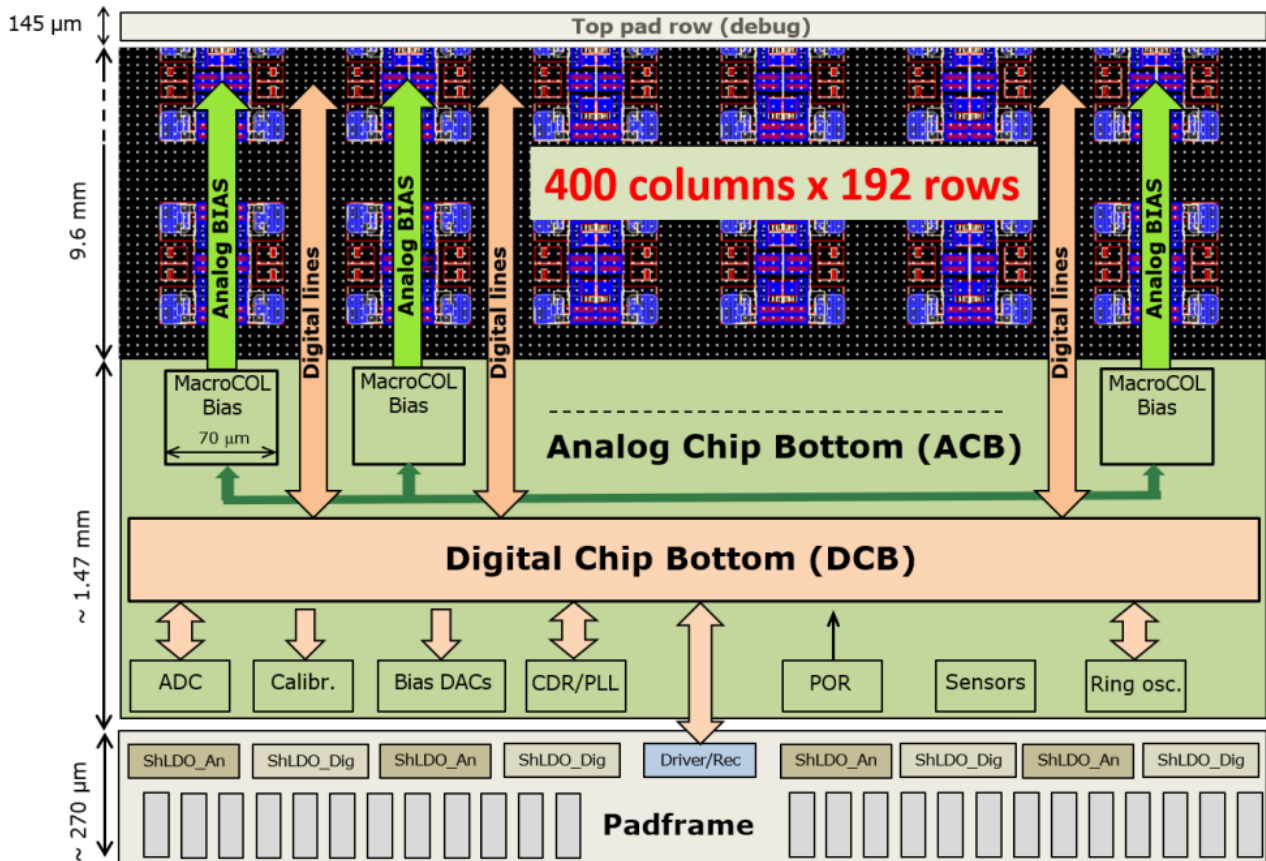


Figura 4.1: Struttura del chip RD53A. Da [42].

lunghi. Una delle due, collocata in figura all'estremità superiore, è provvisoria e sarà rimossa nel chip finale. L'altra, disposta sul lato opposto, contiene i circuiti analogici e digitali necessari a configurare il chip, manipolare i dati in ingresso ed in uscita e a distribuire la tensione di alimentazione. Il fatto che questi siano posizionati su un unico lato rende possibile l'affiancamento di RD53A ad altri 3 chip; tale possibilità sarà sfruttata nei rivelatori a pixel dell'ITk, in cui saranno installati moduli costituiti da 1, 2 o 4 chip connessi ad un unico sensore.

I pixel hanno dimensioni di $50\ \mu\text{m} \times 50\ \mu\text{m}$ e sono organizzati in una matrice di 400 colonne e 192 righe; la regione attiva del chip ha dunque una superficie di $192\ \text{mm}^2$, pari all'83% dell'area totale. Il numero di righe in RD53A è la metà di quello previsto per il chip finale, che avrà pertanto un'altezza maggiore.

Nella regione periferica posta al di sotto della matrice di pixel (con riferimento alla figura 4.1) i circuiti analogici e quelli digitali sono raggruppati rispettivamente nell'*Analog Chip Bottom* (ACB) e nel *Digital Chip Bottom* (DCB). Questa regione ospita, tra le varie componenti, i

registri di configurazione ed i convertitori digitale-analogico (*Digital (to) Analog Converter*, DAC) per tradurre, ove necessario, valori digitali in tensioni o correnti (ad esempio, le tensioni di soglia dei discriminatori, come illustrato nella sezione 4.4.2).

Il chip è montato su una scheda detta *Single Chip Card* (SCC), di dimensioni 10 cm × 10 cm, alla quale è collegato tramite sottili fili metallici, utilizzati per il trasporto dei segnali; tale tecnica è indicata come *wire bonding*. Queste connessioni sono visibili in figura 4.2, nella quale sono rappresentate la SCC ed un ingrandimento del chip.

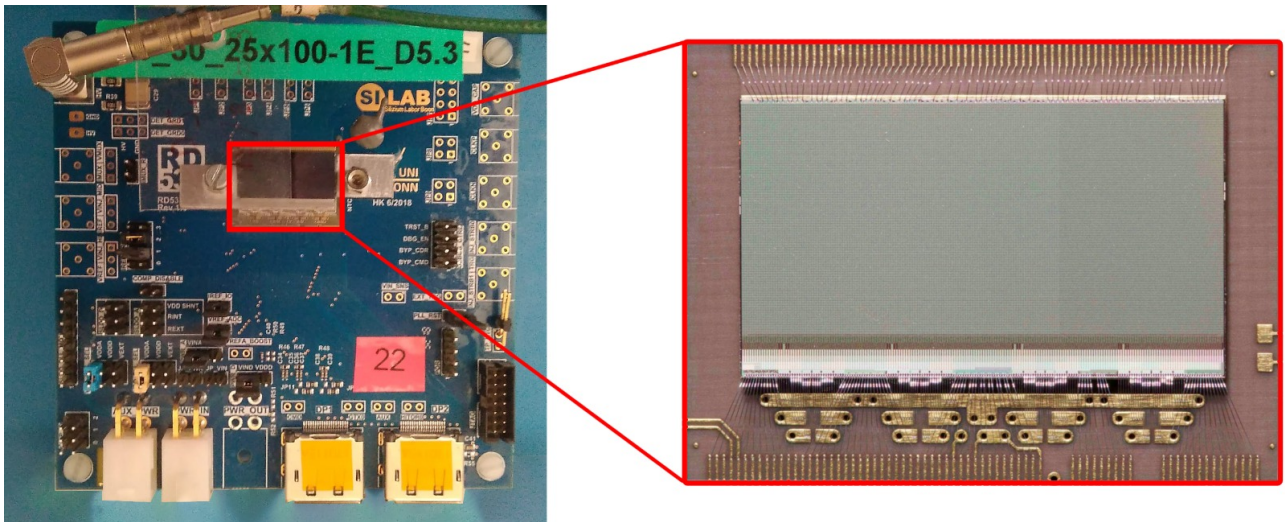


Figura 4.2: *Single Chip Card* (a sinistra) ed ingrandimento del chip RD53A al quale è connesso tramite *wire bonding* (a destra).

4.2.1 La matrice di pixel

Dal punto di vista della logica di lettura ed indirizzamento, i pixel sono organizzati in blocchi di dimensione 8×8 , ciascuno indicato come *core*. I 64 front-end presenti in un core sono raggruppati in 16 blocchi da 4 pixel ciascuno detti *isole analogiche*, disposti sopra un "mare" digitale. Ogni core fornisce i bit di configurazione alle rispettive isole e processa le informazioni da esse ricevute. La disposizione delle isole analogiche è rappresentata in figura 4.3, nella quale è indicata anche la posizione di uno dei pixel del sensore e del bump-bond utilizzato per il suo collegamento.

La disposizione dei bump-bond permette la connessione tra RD53A e sensori con pixel di dimensioni $50 \mu\text{m} \times 50 \mu\text{m}$ (come nel caso ipotizzato in figura) oppure $25 \mu\text{m} \times 100 \mu\text{m}$, corrispondenti alle due ipotesi avanzate per il rivelatore a pixel di ATLAS ed indicate nelle tabelle 3.1 e 3.2; la possibilità di collegare un unico chip a sensori con pixel di diverse dimensioni è stata illustrata nella sezione 3.5.1.

4.3 Alimentazione del chip

Il chip può essere alimentato in modalità passiva o attiva. Nel primo caso i circuiti analogici e quelli digitali sono alimentati direttamente dal generatore; in modalità attiva, invece, la tensione è fornita attraverso dei regolatori, detti *Shunt - LDO regulator* (ShuLDO), che filtrano la

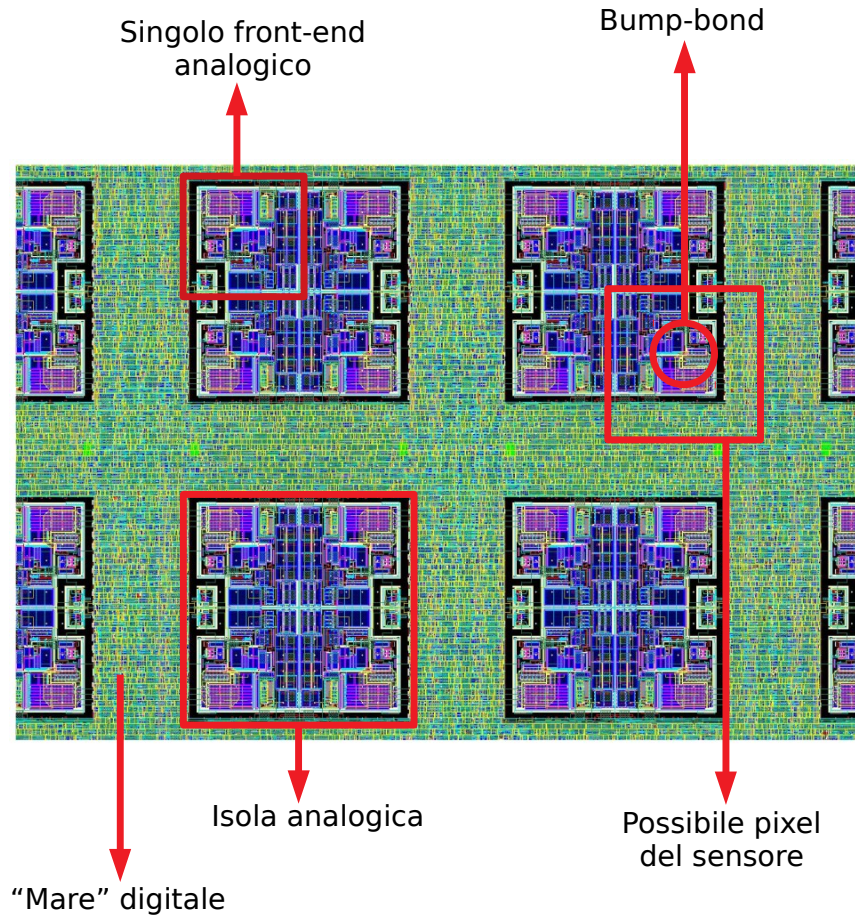


Figura 4.3: Disposizione delle isole analogiche disposte sopra un "mare" digitale. È indicata anche la possibile posizione di uno dei pixel del sensore e del bump-bond utilizzato per il suo collegamento.

tensione ricevuta dall'alimentatore esterno. Un regolatore è un dispositivo utilizzato per mantenere costante un livello di tensione; il termine "shunt" (letteralmente, "deviare") deriva dal fatto che tale obiettivo è raggiunto deviando verso massa un eventuale eccesso di corrente, mentre LDO suggerisce che il dispositivo lavora a basso *dropout* (*Low-DropOut*), ovvero con una piccola differenza tra la tensione in ingresso e quella in uscita. RD53A è dotato di due regolatori ShuLDO per poter alimentare separatamente la componente analogica (SLDO-A) e quella digitale (SLDO-D) a partire dalla tensione fornita da un unico alimentatore. A differenza del caso passivo, in cui l'alimentazione dei front-end può essere controllata agendo direttamente sul generatore, in modalità attiva è possibile cambiare la tensione in uscita dai regolatori modificando i parametri di configurazione del chip, contenuti in un apposito file (si veda la sezione 5.4).

La SCC e RD53A sono stati progettati in modo da consentire un'alimentazione di tipo seriale, come rappresentato in figura 4.4. Nell'Inner Detector attualmente installato ad ATLAS i moduli sono alimentati in modo parallelo da un generatore posizionato all'esterno del volume attivo del rivelatore. Ciò comporta l'utilizzo di cavi di alimentazione di elevata lunghezza e la necessità di fornire, nell'ipotesi ideale di dissipazione nulla, una corrente pari a nI , in cui n rappresenta il numero di moduli collegati in parallelo tra loro e I la corrente necessaria al loro funzionamento. Un'alimentazione di tipo seriale, invece, permette di ridurre la lunghezza dei cavi e dunque di

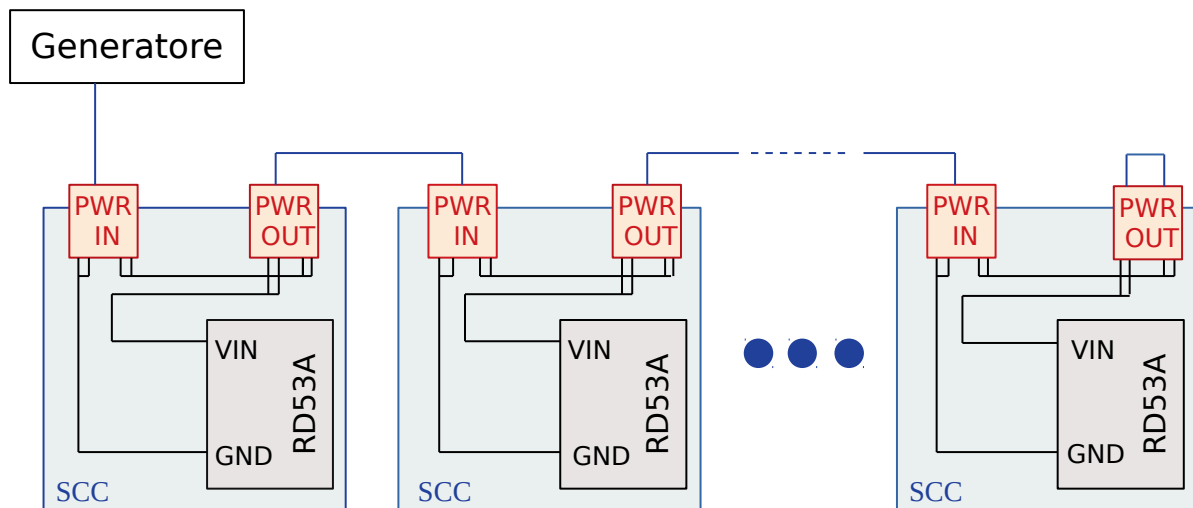


Figura 4.4: Alimentazione seriale dei chip RD53A.

minimizzare il materiale con cui le particelle che si propagano nel rivelatore devono interagire e la dispersione del segnale; inoltre, la corrente necessaria all'alimentazione dello stesso numero di moduli si riduce di un fattore n . Tali motivazioni hanno spinto la collaborazione ATLAS a valutare l'efficienza del sistema seriale; l'esito favorevole degli studi condotti ha guidato la realizzazione di RD53A.

4.4 Elettronica di front-end di un generico chip di lettura

L'elettronica di front-end implementata in ciascun pixel provvede all'amplificazione del segnale, alla sua discriminazione dal rumore, alla sua digitalizzazione ed al salvataggio dell'informazione ottenuta. Tali attività sono svolte da due diverse componenti, una analogica ed una digitale. La più generica logica di front-end, sulla base della quale si sviluppano quelle utilizzate da tutti i chip di lettura, è rappresentata in figura 4.5; le quantità ed i segnali gestibili dall'utente sono divisi tra quelli globali, relativi all'intera matrice ed indicati in blu, e quelli locali, associati ai singoli pixel ed indicati in rosso. Nelle sezioni successive è descritto il funzionamento della componente analogica e di quella digitale del circuito di front-end.

4.4.1 Componente analogica

La componente analogica riceve il segnale dal sensore oppure da un circuito di iniezione, il quale simula il rilascio di energia dovuto all'interazione di una particella con il sensore. Questo permette di verificare il corretto funzionamento del front-end durante le fasi di ottimizzazione del chip, nota la carica iniettata. Il circuito di iniezione può essere abilitato globalmente, ovvero all'interno di tutta la matrice, oppure solo in alcuni pixel, in modo da selezionare quelli di interesse.

Il segnale di ingresso è ricevuto da un amplificatore operazionale con una capacità di feedback C_f ; tale dispositivo è in genere un *Charge Sensitive Amplifier* (CSA), ovvero un elemento che integra la corrente che fluisce attraverso il terminale invertente generando una tensione

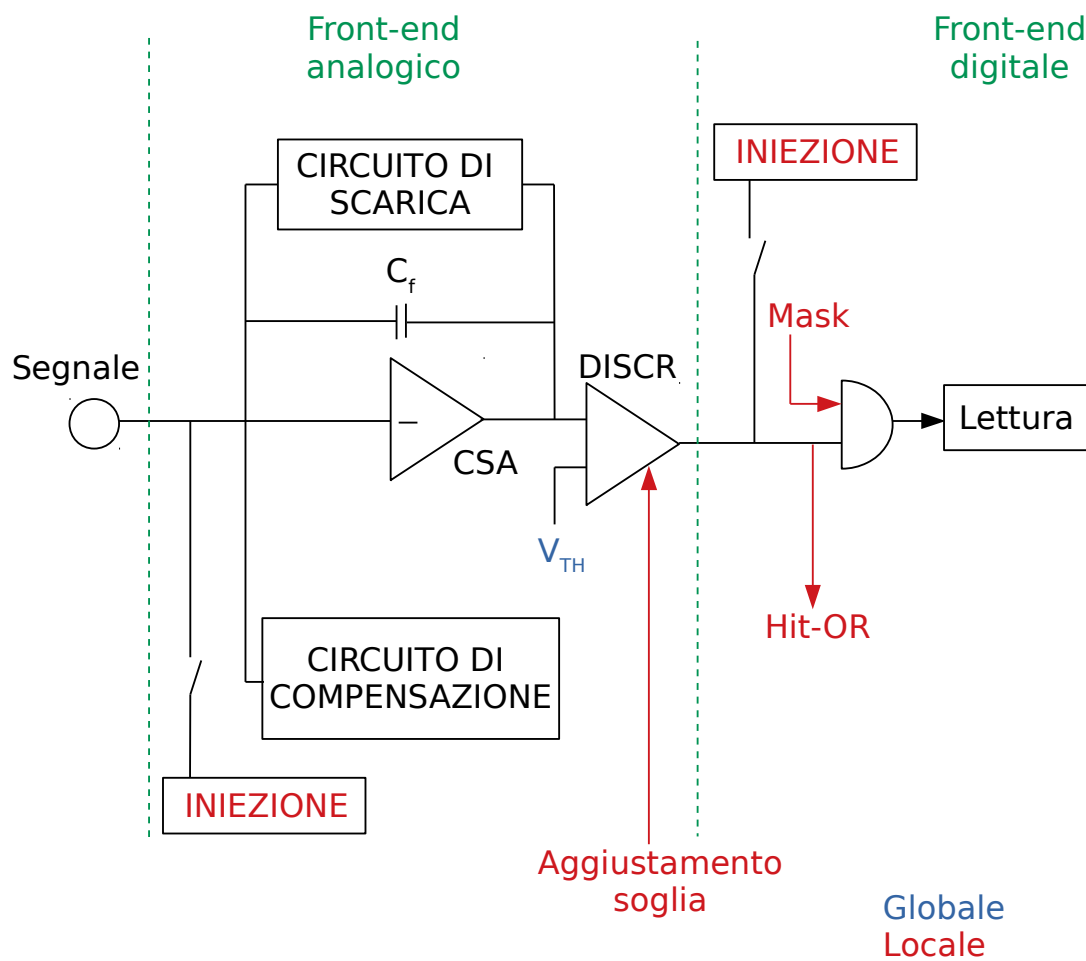


Figura 4.5: Schematizzazione di una generica logica di front-end utilizzata dai chip di lettura di pixel.

di ampiezza proporzionale alla carica in ingresso. In figura 4.6 è rappresentato un esempio di due segnali di diversa ampiezza all'ingresso del CSA (a sinistra) ed i relativi segnali in uscita (al centro). Quando tutta la carica è stata raccolta l'uscita dell'amplificatore diventa stabile e resta in tale condizione fino a che non si verifica la scarica del condensatore; nel caso in cui si sia in saturazione ciò implica un tempo morto, in quanto il passaggio di una nuova particella produrrebbe un segnale che si sovrapporrebbe al precedente, rendendoli indistinguibili. Per questo motivo è necessario poter indurre la scarica di C_f ; in tal modo il segnale in uscita dal CSA ritorna al valore assunto in assenza di carica al suo ingresso, indicato come *baseline*. La soluzione più semplice è rappresentata dall'introduzione di una resistenza R in parallelo a C_f , la quale consente di scaricare il condensatore dopo un tempo $\tau = RC_f$. Considerando un tipico valore di C_f utilizzato nei CSA, pari a $10 \mu\text{F}$, per poter ottenere $\tau = 1 \mu\text{s}$ è necessario utilizzare $R = 100 \text{ M}\Omega$; tuttavia, è difficile integrare resistenze così elevate all'interno di un pixel. Per questo motivo nella maggior parte dei front-end si opta per una soluzione alternativa e si implementano dei circuiti che forniscono una corrente che induce la scarica del condensatore, ottenendo in uscita dall'amplificatore un segnale di durata proporzionale alla carica in ingresso; l'azione di questo circuito è rappresentata a destra in figura 4.6. Modificando il valore della corrente in uscita da

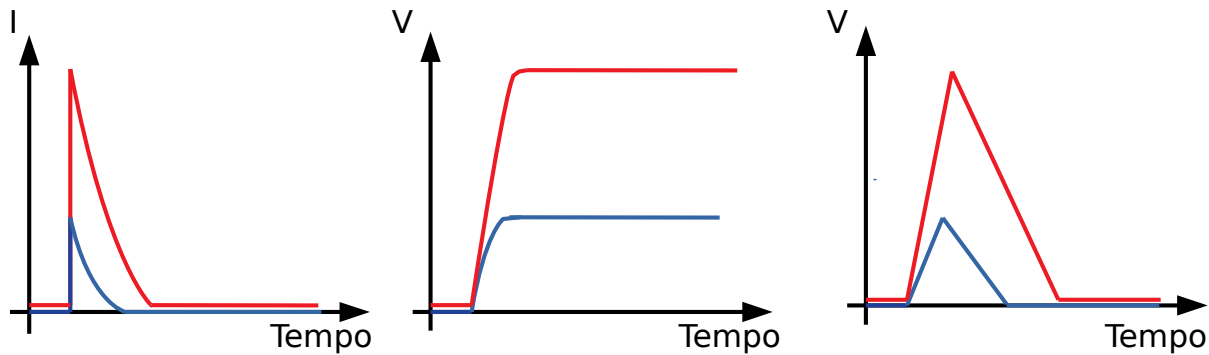


Figura 4.6: Sinistra: esempio di due segnali di diversa ampiezza all'ingresso del CSA. Centro: segnale in uscita dal CSA corrispondente ai due ingressi rappresentati a sinistra. Destra: segnale in uscita dal CSA corrispondente ai due ingressi rappresentati a sinistra nel caso in cui si adoperi un circuito di scarica del condensatore.

tale circuito, è possibile aumentare o diminuire la velocità di scarica del condensatore, variando la proporzionalità tra la quantità di carica all'ingresso del CSA e la durata del segnale in uscita. Tale possibilità è illustrata in figura 4.7, nella quale i diversi segnali corrispondono allo stesso ingresso e I_{scarica} indica la corrente fornita dal circuito di scarica.

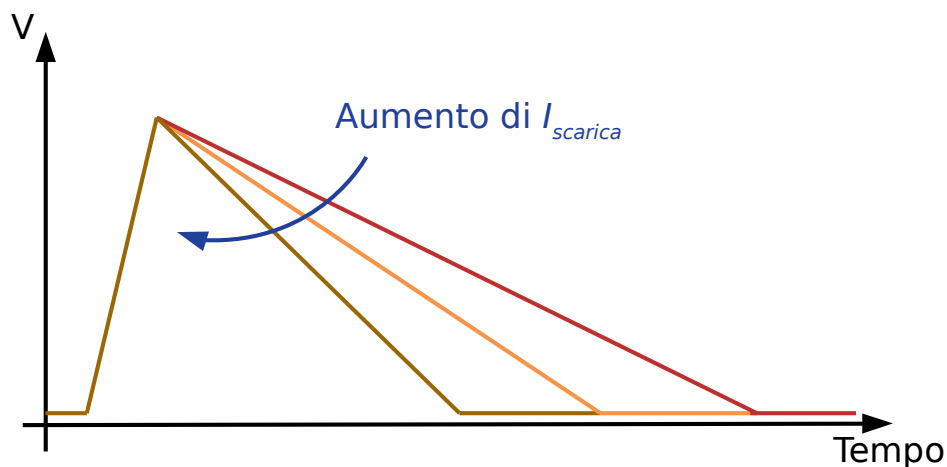


Figura 4.7: Influenza della corrente fornita dal circuito di scarica del condensatore (I_{scarica}) sulla velocità di ritorno alla baseline del segnale in uscita dall'amplificatore.

Un secondo circuito che può essere abbinato al CSA è quello addetto alla compensazione della corrente di leakage, il quale devia tale corrente evitando che raggiunga l'ingresso dell'amplificatore e venga integrata insieme al segnale.

4.4.2 Componente digitale

La tensione generata dall'amplificatore è inviata ad un discriminatore (DISCR in figura 4.5), il quale la confronta con un valore di riferimento, indicato come *soglia* (V_{TH} in figura 4.5). Nel caso in cui l'uscita del CSA sia maggiore rispetto alla soglia il circuito produce un *hit*, ovvero un

segnale in uscita. Ciò può accadere per effetto di un segnale in ingresso prodotto al passaggio di una particella o, equivalentemente, simulato tramite il circuito di iniezione o generato dal rumore. Per ottimizzare le prestazioni del front-end, la progettazione dei pixel avviene in modo che le fluttuazioni da rumore elettronico abbiano un'ampiezza che mediamente è sensibilmente inferiore all'ampiezza media di un segnale prodotto da una particella. Per minimizzare il rumore è dunque preferibile avere soglie alte. Tuttavia, per massimizzare l'efficienza di rivelazione garantendo di rimanere sensibili a segnali che fluttuano al di sotto del valore più probabile è preferibile mantenere la soglia più bassa possibile. Pertanto, la scelta della tensione di riferimento deve derivare da un compromesso tra le due esigenze.

Il valore della soglia è impostato globalmente per tutti i pixel; tuttavia, a causa di variazioni delle caratteristiche dei transistor con cui sono realizzati i circuiti tale riferimento non è uguale per tutti i pixel. Queste variazioni, dovute a fluttuazioni nel processo di produzione, sono indicate come *mismatch* dei transistor. Un metodo per ridurre la dispersione di soglia dovuta a tale mismatch è rappresentato dall'utilizzo di un DAC che fornisca una tensione adeguata a compensarla. Poiché questa dispersione non è costante in ogni pixel, ognuno di essi è associato ad un registro di compensazione che determina il valore locale di tale tensione aggiuntiva.

Un secondo circuito di iniezione è implementato all'uscita del discriminatore ed è adoperato per inviare un impulso digitale che simuli la ricezione di un segnale sopra soglia. Tale possibilità è utilizzata durante le fasi di ottimizzazione del chip per verificare il corretto funzionamento della logica di lettura dei segnali; analogamente al precedente circuito di iniezione, può essere abilitato globalmente o solo per determinati pixel.

Il discriminatore fornisce un segnale di livello alto qualora la tensione in ingresso sia superiore rispetto alla soglia. La durata di tale segnale, indicata come *Time Over Threshold* (ToT), è pertanto proporzionale alla carica generata all'ingresso dell'amplificatore, come illustrato a sinistra in figura 4.8 (si ricorda che la larghezza del segnale in uscita dal CSA è proporzionale all'ampiezza del segnale al suo ingresso). I segnali in uscita dal discriminatore corrispondenti a quelli in figura 4.7 sono illustrati a destra in figura 4.8, nella quale è visibile l'influenza della corrente $I_{scarica}$ sul valore di ToT corrispondente alla stessa quantità di carica all'ingresso del front-end.

4.4.2.1 Lettura delle informazioni

La lettura delle informazioni ricavate da ogni pixel prevede il conteggio del ToT, il suo salvataggio ed il suo trasferimento al mondo esterno. Tali operazioni possono essere inibite nel caso in cui un pixel manifesti un comportamento scorretto; il segnale di inibizione è indicato in figura come *Mask*. La lettura di un pixel è abilitata dalla ricezione di un segnale, detto *trigger*; la modalità di generazione del trigger varia in base all'applicazione per la quale un chip è progettato. Se la frequenza dei segnali è talmente bassa da consentire il loro salvataggio sul disco senza prima operare alcuna selezione è possibile generare il trigger sfruttando l'informazione fornita dal discriminatore: nel caso in cui la sua uscita sia alta è prodotto un segnale, detto *hit-OR*, che abilita la lettura del pixel, fungendo da trigger. In questo modo tutti gli hit verificatisi entro un certo intervallo di tempo (predefinito dall'utente) sono letti dai circuiti digitali. In esperimenti nei quali la frequenza dei segnali è molto alta, quale è ATLAS, è necessario individuare gli hit di effettivo interesse tra tutti quelli ricevuti (ovvero quelli associati ad eventi in cui la logica di trigger dell'esperimento ha riconosciuto la probabile produzione di un evento interessante); solo questi saranno poi processati dal sistema di acquisizione dati e salvati su disco. Ciò è possibile

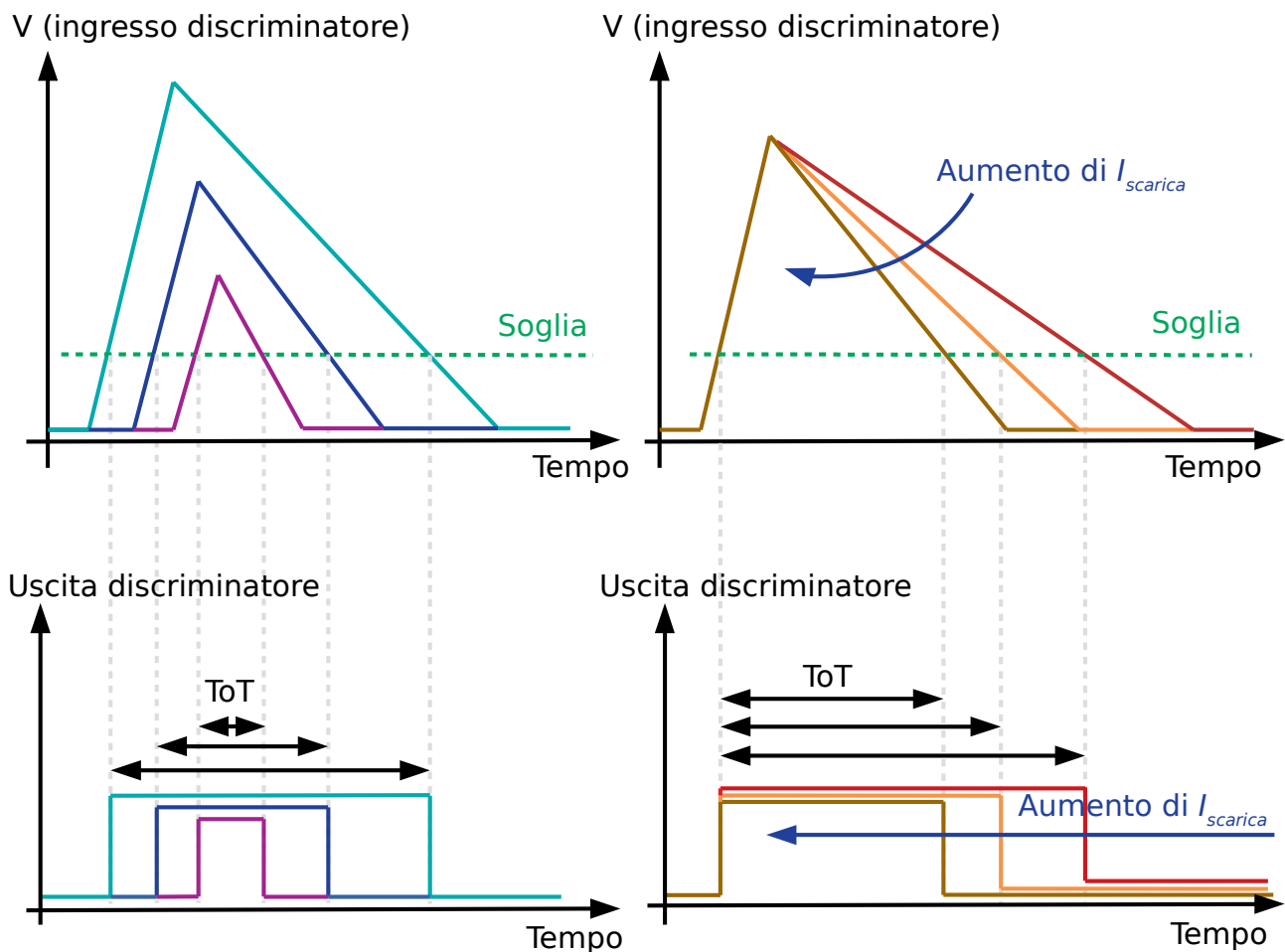


Figura 4.8: Sinistra: relazione tra i segnali in ingresso (in alto) ed in uscita (in basso) da un discriminatore al variare della quantità di carica all'ingresso del front-end. Destra: relazione tra i segnali in ingresso (in alto) ed in uscita (in basso) da un discriminatore a parità di carica all'ingresso del front-end al variare della corrente fornita dal circuito di scarica del condensatore di feedback del CSA.

utilizzando un segnale di trigger esterno al front-end, che abilita la lettura di un determinato sottoinsieme di hit sulla base di una determinata logica.

4.5 Front-end analogico di RD53A

In RD53A sono implementate tre logiche di front-end, indicate come *lineare*, *differenziale* e *sincrona*, che rappresentano diverse varianti di quella illustrata in figura 4.5. In figura 4.9 è rappresentata la matrice dei pixel di RD53A con l'indicazione delle aree in cui sono implementate le tre tipologie di front-end.

Il front-end sincrono è utilizzato per la lettura di 16 colonne di core, mentre gli altri due sono implementati in 17 colonne di core ciascuno; questa scelta è dettata dal fatto che la logica sincrona dissipa una quantità di potenza leggermente superiore rispetto alle altre.

Le caratteristiche principali dei diversi front-end sono descritte nei paragrafi successivi; per semplicità di rappresentazione saranno sottointesi i circuiti di iniezione analogica e digitale ed i

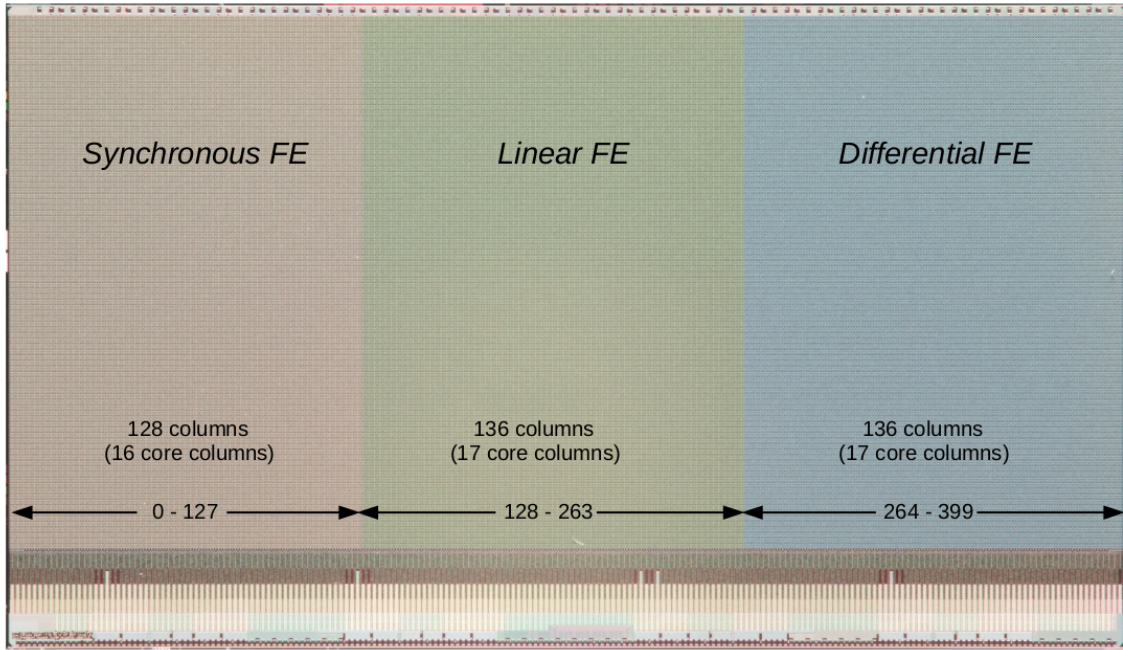


Figura 4.9: Disposizione sulla matrice di pixel di RD53A dei diversi front-end implementati. Da [42].

segnali di Hit-Or e di Mask illustrati in figura 4.5, in quanto comuni alle tre logiche.

4.5.1 Front-end lineare

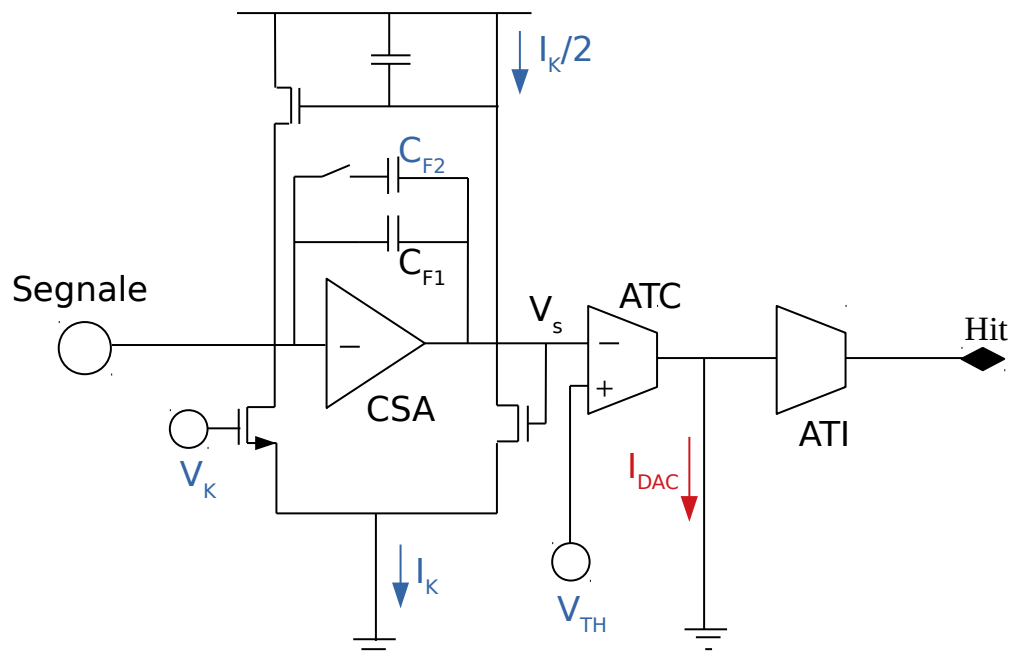
Il front-end lineare è rappresentato in figura 4.10. Al condensatore di feedback C_{F1} dell'amplificatore di carica è collegata in parallelo una seconda capacità, C_{F2} , che può essere connessa o disconnessa dal circuito a seconda che si voglia ottenere un guadagno di amplificazione rispettivamente minore o maggiore.

La scarica del condensatore e la compensazione della corrente di leakage sono ottenute utilizzando un unico circuito, detto di Krummenacher [43], alimentato dalla tensione V_K . Il valore della corrente che fluisce in questo circuito, I_K , può essere impostato globalmente modificando i parametri di configurazione che agiscono sul DAC da cui è generata V_K ; il suo effetto è quello rappresentato in figura 4.8.

Il confronto tra il segnale in uscita dall'amplificatore e la tensione di riferimento V_{TH} avviene in due passaggi successivi, che costituiscono una variante del discriminatore illustrato in figura 4.5. Il primo stadio è rappresentato da un amplificatore a transconduttanza (ATC in figura), il quale fornisce in uscita una corrente I_{out} proporzionale alla differenza di tensione tra gli ingressi:

$$I_{out} = (V_S - V_{TH}) \cdot g$$

in cui g indica la transconduttanza del dispositivo. Tale corrente, pertanto, ha segno positivo o negativo in base alla relazione tra V_S e V_{TH} . Il secondo stadio di discriminazione è costituito da un amplificatore a transimpedenza (ATI in figura), il quale converte la corrente ricevuta in ingresso in tensione. Il segnale in uscita, pertanto, sarà ad un livello alto o basso a seconda del fatto che quello fornito dal CSA sia maggiore o minore rispetto alla soglia.



Globale
Locale

Figura 4.10: Schematizzazione del front-end lineare.

A differenza di quanto illustrato in figura 4.5, il confronto con il segnale di riferimento avviene in corrente e non in tensione; di conseguenza, la dispersione della soglia è ridotta utilizzando un DAC che genera una corrente, indicata in figura come I_{DAC} . Il suo valore è regolato agendo su 4 bit.

4.5.2 Front-end differenziale

Il front-end differenziale è rappresentato in figura 4.11. In parallelo alla capacità di feedback C_{F1} è inserito il condensatore C_{F2} , il quale può essere connesso al circuito riducendo il guadagno di un fattore 1.6. La corrente che induce la scarica dei condensatori è indicata in figura come I_{ff} , il circuito di compensazione della corrente di leakage come LCC (*Leakage Current Compensation*). Quest'ultimo può essere disabilitato per ridurre il rumore all'interno del front-end nel caso in cui la corrente di leakage abbia un valore talmente basso da non compromettere la misura. Tale possibilità non è prevista nel caso lineare, in cui la compensazione è fornita dal circuito di Krummenacher che non può essere disabilitato in quanto responsabile anche della scarica del condensatore di feedback.

L'uscita dell'amplificatore ed il suo ingresso, virtualmente a massa, sono inviati allo stadio successivo, indicato come *precomparatore* (*PreComp*) in figura; esso riceve in ingresso anche due valori di soglia, indicati in figura come V_{TH1} e V_{TH2} . Nel precomparatore sono formati due segnali di opposta polarità, rappresentati dalla differenza tra l'ingresso e l'uscita del CSA; questi

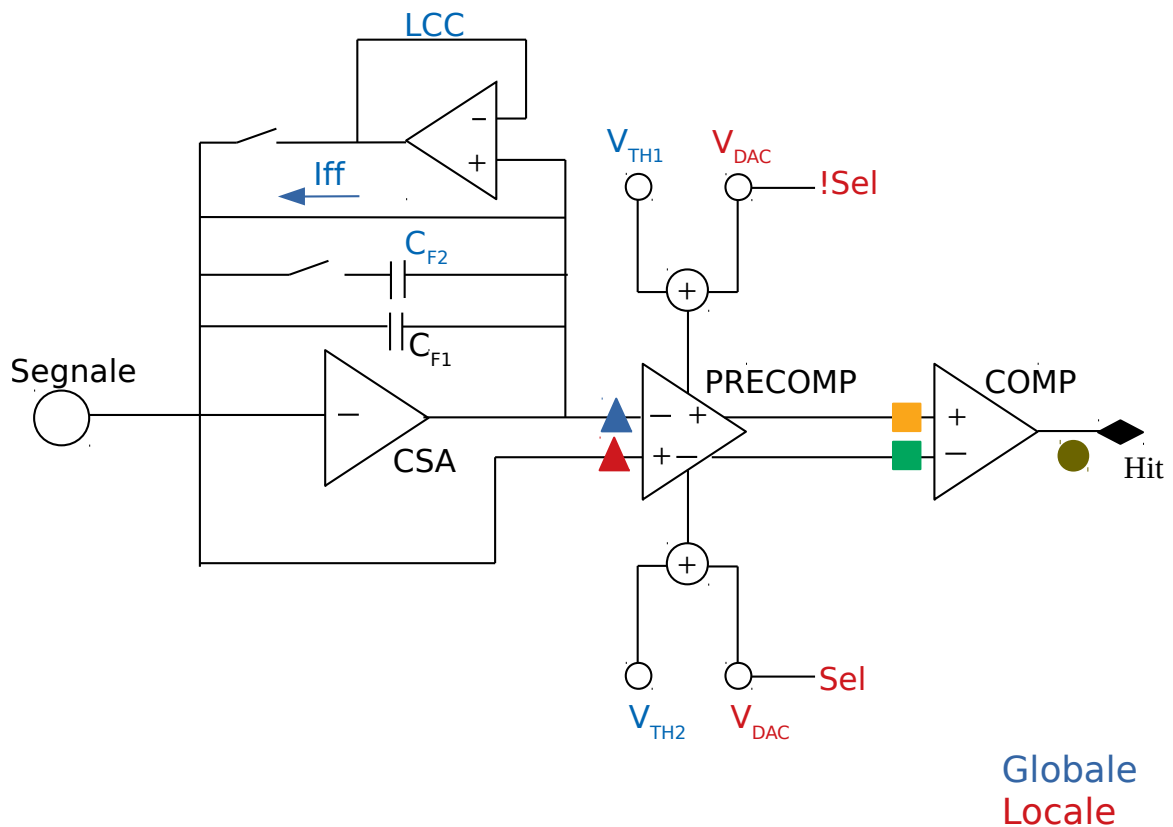


Figura 4.11: Schematizzazione del front-end differenziale.

sono sommati ai valori di riferimento, i quali dunque sbilanciano le due uscite del precomparatore, aumentando o diminuendo la loro distanza. La dispersione di soglia è ridotta grazie ad un DAC a 5 bit; di questi, 4 sono utilizzati per la generazione della tensione, il quinto (indicato in figura come *Sel*) seleziona la soglia sulla quale agire.

Le uscite del precomparatore sono inviate ad un comparatore (*Comp* in figura); questo fornisce in uscita un segnale binario il cui stato cambia nell'istante in cui avviene un'intersezione tra i segnali in ingresso. Aumentare la distanza tra le uscite del precomparatore corrisponde quindi ad incrementare l'intensità che deve avere il segnale rivelato dal CSA affinché gli ingressi del comparatore si intersechino. Ciò equivale ad aumentare la tensione di soglia, che è pertanto definita in modo differenziale.

In figura 4.12 è rappresentato un esempio di come un generico segnale di ingresso è processato dal front-end differenziale; la legenda utilizzata indica la posizione dei diversi segnali nel circuito in figura 4.11.

4.5.3 Front-end sincrono

Il front-end sincrono è rappresentato in figura 4.13. Il primo stadio è costituito da un amplificatore di carica collegato ad un circuito di Krummenacher. A differenza di ciò che accade negli altri circuiti, entrambi i condensatori di feedback possono essere collegati o meno al CSA,

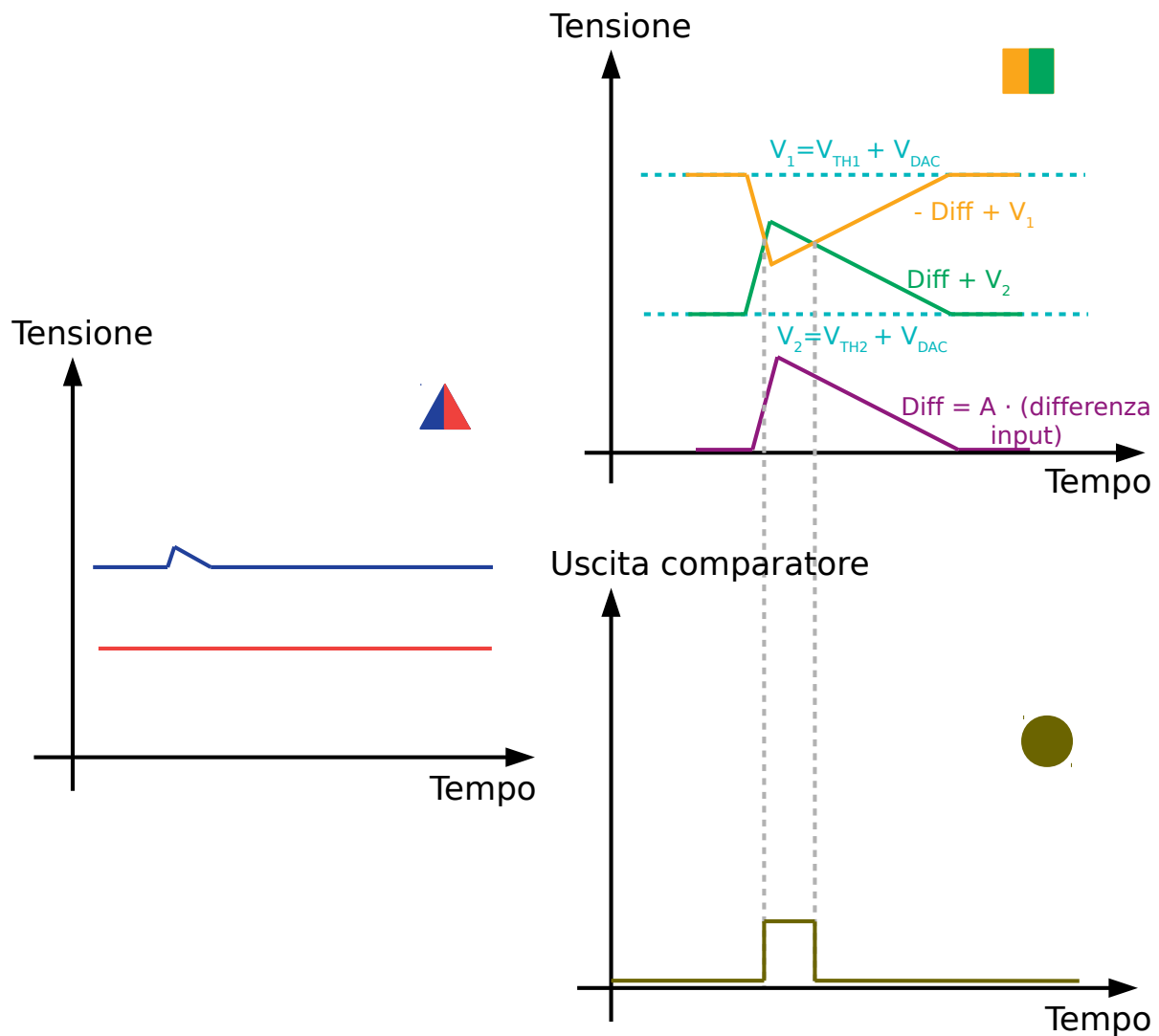


Figura 4.12: Esempio di segnali in uscita dal CSA (a sinistra), dal precomparatore (destra, in alto) e dal comparatore (destra, in basso) nel front-end differenziale.

permettendo dunque di scegliere tra tre possibili guadagni: massimo, connettendo C_{F1} (di valore pari a 1.6 fF), intermedio, connettendo C_{F2} (di valore pari a 2.4 fF) e minimo, connettendo entrambi i condensatori (ottenendo quindi una capacità di 4 fF).

Il segnale in uscita dal CSA è inviato ad uno dei due ingressi di un amplificatore differenziale (AD in figura); il secondo ingresso è utilizzato per fornire la tensione di soglia, V_{TH} . L'accoppiamento tra i due stadi è effettuato in modo capacitivo (tramite il condensatore C_{AC}), eliminando così il problema della dispersione della baseline dovuto al mismatch dei transistor con cui è realizzato il CSA. Affinché i due ingressi dell'AD siano riferiti alla stessa massa, a valle del condensatore C_{AC} è posto un generatore di tensione che introduce una nuova baseline, indicata in figura come V_{BL} . Per effetto del mismatch dei transistor con cui è realizzato l'amplificatore differenziale, anche nel caso in cui i due ingressi siano idealmente allo stesso potenziale esiste una differenza di tensione tra di essi, indicata in figura come V_{OS} . Tale valore può essere diverso

in ogni pixel e ciò determina una dispersione della soglia; per ridurla si utilizza un meccanismo, detto di *auto-zero*, che campiona la differenza di potenziale tra i due ingressi di AD e la sottrae alla misura.

$$A \cdot (V_S + V_{BL} + V_{OS} - V_{TH}) - A \cdot V_{OS} = A \cdot (V_S + V_{BL} - V_{TH}).$$

I due segnali in uscita dall'amplificatore differenziale sono inviati all'ingresso di un latch, il quale memorizza l'eventuale presenza di un hit ed è letto utilizzando un segnale indicato come *strobe*.

4.6 Front-end digitale di RD53A

Come anticipato nella sezione 4.4 la lettura dei pixel di RD53A, effettuata dalla componente digitale dell'elettronica di front-end, è abilitata dalla ricezione di un segnale di trigger, che individua tra tutti gli hit verificatisi quelli di effettivo interesse. La selezione degli eventi ad ATLAS è effettuata sulla base delle informazioni provenienti da diversi sotto-rivelatori ed esiste un certo ritardo tra l'istante in cui ognuno di essi rivela un segnale e quello in cui riceve il trigger che abilita la lettura; tale differenza temporale è indicata come *latenza*. Ciò comporta la necessità di salvare localmente i dati per un certo intervallo, al termine del quale sono inviati al livello di processamento successivo oppure eliminati a seconda del fatto che si sia ricevuto o meno un segnale di trigger.

In RD53A sono adottate due diverse logiche di lettura, denominate *Distributed Buffer Architecture* (DBA) e *Central Buffer Architecture* (CBA). La prima è utilizzata dai front-end lineare e differenziale, la seconda da quello sincrono. In entrambi i casi esiste un'area condivisa da un certo numero di pixel, indicata come *regione*, contenente le informazioni temporali riguardo agli hit registrati. Le differenze tra le due architetture riguardano le dimensioni di tale regione ed il modo in cui sono salvati i valori di ToT.

La formattazione dei dati letti dalle due diverse logiche è la stessa, in modo da agevolare il loro successivo processamento in una logica comune.

4.6.1 Distributed Buffer Architecture

La logica DBA è schematizzata in figura 4.14. Ogni regione, rappresentata in figura, contiene 4 pixel, affiancati lungo un'unica riga. Ognuno di essi è dotato di un contatore di ToT a 4 bit, il quale funziona in maniera sincrona con un clock a 40 MHz, corrispondente alla frequenza delle interazioni dei fasci ad LHC; il ToT è pertanto calcolato in unità di bc (*bunch crossing*). Avendo a disposizione 4 bit, il più alto valore di ToT registrabile è pari a 15 bc; nel caso in cui l'uscita del discriminatore nel front-end analogico sia ancora alta quando il ToT ha raggiunto il valore di 15 bc viene salvata tale quantità e si azzerà il contatore. I valori misurati sono salvati all'interno degli 8 registri presenti in ciascun pixel.

Nell'area comune a tutta la regione sono contenuti 8 registri da 9 bit ciascuno, ognuno collegato ad un registro dedicato al salvataggio del ToT per ciascuno dei 4 pixel. La rivelazione di un hit in un pixel attiva un contatore in uno dei registri nell'area comune. I ToT sono conservati in memoria fino a che tale contatore non raggiunge il valore massimo, misurato anch'esso in termini di bc e pari a 511 bc; poiché una frequenza di 40 MHz corrisponde ad un intervallo temporale di 25 ns, il massimo tempo di latenza permesso per la ricezione del segnale di trigger è $511 \times 25 \text{ ns} = 12.8 \text{ }\mu\text{s}$. Quando il contatore di latenza ha raggiunto tale valore i ToT sono eliminati qualora non si sia ricevuto il trigger; in caso contrario viene abilitata la lettura dei ToT salvati nella regione. In questo modo saranno letti anche i valori nulli; la logica DBA, pertanto, è indicata nel caso in cui la frequenza di registrazione degli hit sia molto alta, in modo che la quantità di memoria sprecata per salvare ToT nulli, e quindi non significativi, sia limitata. Per questo motivo all'interno di una regione sono raggruppati soltanto 4 pixel.

4.6.2 Central Buffer Architecture

La logica CBA può essere schematizzata come in figura 4.15. Ogni regione, rappresentata in

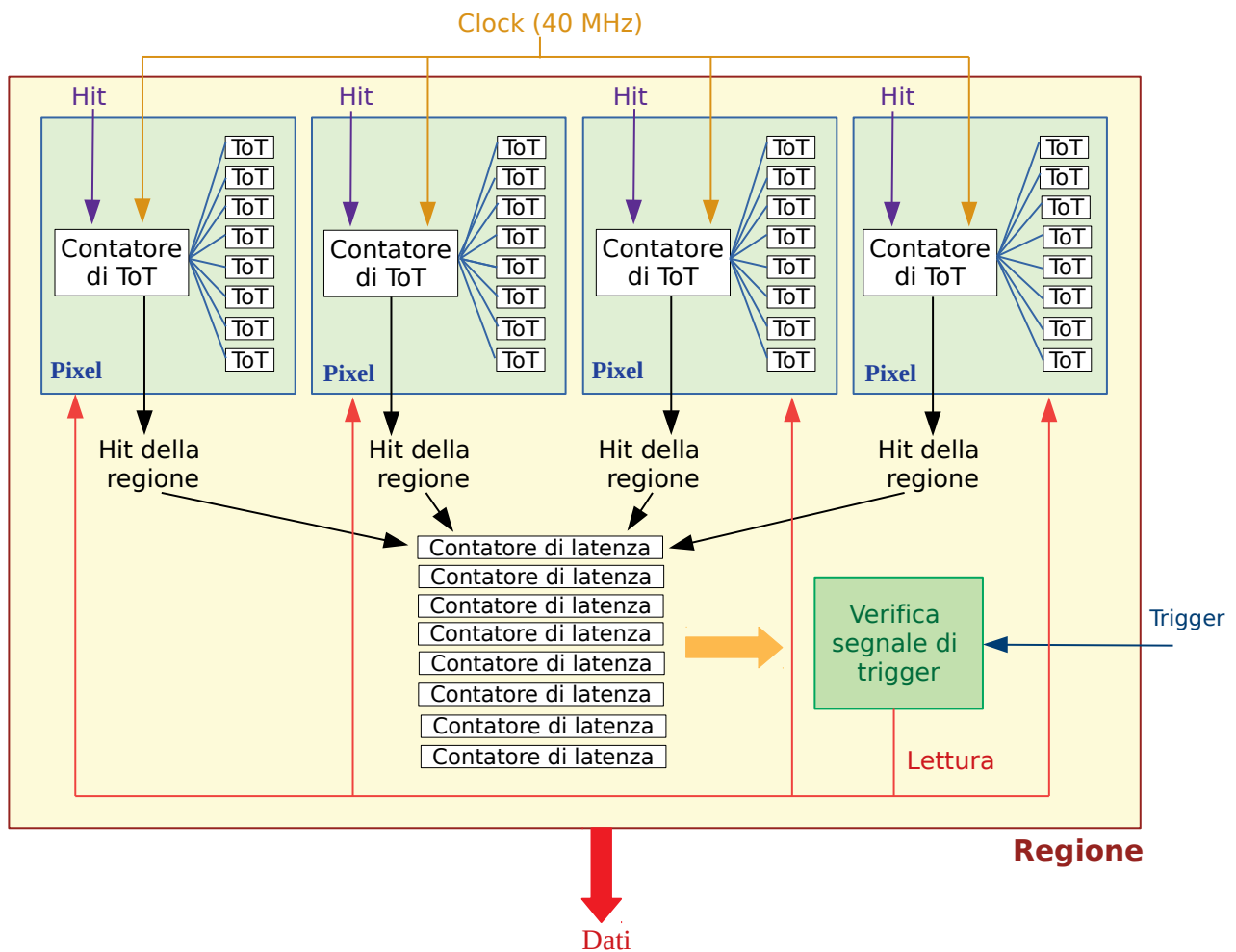


Figura 4.14: Schematizzazione della logica DBA.

figura, contiene 16 pixel, disposti in una matrice 4×4 .

La principale differenza rispetto alla logica DBA è rappresentata dal fatto che i valori di ToT non sono salvati in ogni pixel ma nell'area comune a tutta la regione. Ciò consente di memorizzare soltanto i valori non nulli, risparmiando quindi memoria. Tuttavia, è necessario creare una mappa degli hit registrati, detta *hitmap*.

La hitmap è salvata all'interno di un buffer per il tempo necessario al completamento del calcolo dei ToT; tale buffer contiene anche 16 registri da 9 bit ciascuno, utilizzati per il conteggio del tempo di latenza. L'aumento del numero di registri rispetto all'architettura precedente riflette l'incremento del numero di pixel all'interno di una regione. Ogni mappa è creata dalla ricezione di un hit; il buffer consente di salvare contemporaneamente 4 hitmap.

I ToT sono inviati allo stadio successivo, indicato in figura come *zero-suppression*, in cui avviene l'eliminazione dei valori ToT = 0 bc.

La mappa degli hit, il contenuto dei contatori di latenza e i ToT non nulli sono inviati in un secondo buffer, il quale verifica l'associazione tra gli hit registrati ed il segnale di trigger. Nel buffer possono essere salvati 8 ToT; è possibile pertanto che alcuni valori siano persi. I corrispondenti hit, tuttavia, sono comunque registrati all'interno della mappa.

La logica CBA è particolarmente utile nel caso in cui la frequenza degli hit attesa sia bassa e

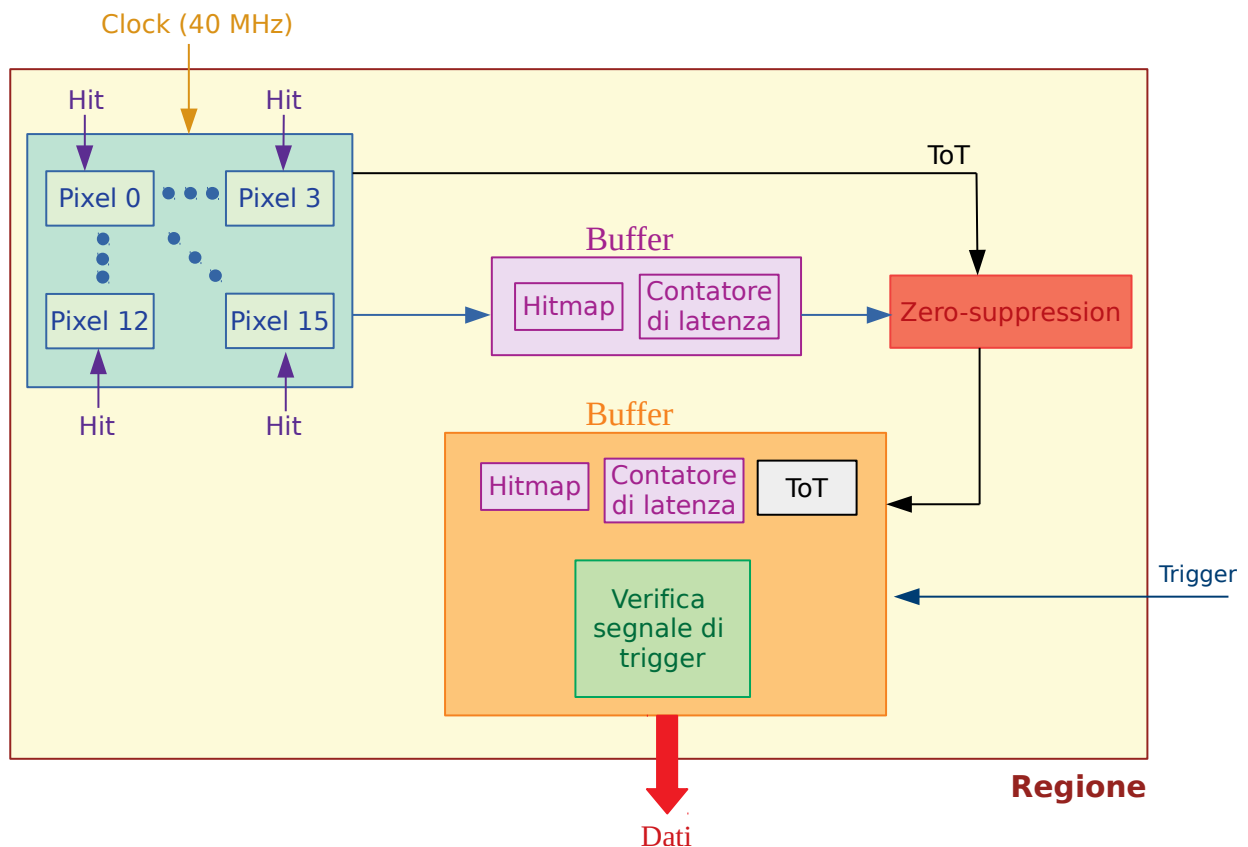


Figura 4.15: Schematizzazione della logica CBA.

ci siano dunque molti valori di ToT nulli che, se salvati, comporterebbero uno spreco di memoria. Per questo motivo la regione scelta per questa architettura ha estensione maggiore rispetto a quella utilizzata nella DBA.

4.7 Periferia del chip

La periferia del chip contiene le pad utilizzate per il collegamento tramite wire bonding alla Single Chip Card ed i blocchi ACB e DCB.

4.7.1 Analog Chip Bottom

Uno dei principali compiti svolti dall'ACB consiste nella generazione della tensione di alimentazione dei circuiti di front-end, ottenuta tramite dei DAC a 10 bit. La tensione fornita da tali DAC è inviata a 200 diversi circuiti, detti *MacroColumnBias*, i quali alimentano in parallelo coppie di colonne adiacenti della matrice di pixel. In questo modo è possibile limitare ad un'unica coppia di colonne un eventuale danneggiamento di una delle linee di alimentazione. La tensione generata dai DAC ed altri valori analogici, quali quelli di temperatura, possono essere monitorati tramite un convertitore analogico-digitale (*Analog (to) Digital Converter*, ADC) a 12 bit contenuto nell'ACB.

Questo blocco, inoltre, include un *Clock Data Recovery* (CDR), il quale genera tre diversi segnali di clock: uno a 160 MHz, detto *command clock*, uno a 1.28 GHz ed uno a 640 MHz. Quest'ultimo è utilizzato per sincronizzare il command clock del chip con il clock di LHC (a 40 MHz).

Nell'ACB si trovano anche quattro serializzatori, i quali trasmettono serialmente su quattro diverse linee i dati in uscita dal chip, in maniera sincrona con il clock a 1.28 GHz generato dal CDR. Questo trasferimento avviene con una larghezza di banda di 5.12 Gb/s, pari al valore richiesto per i chip che saranno utilizzati nello strato più interno del rivelatore a pixel di ITk. La larghezza di banda dei chip utilizzati nell'IBL è di 160 Mb/s [30]; risulta quindi evidente come RD53A sia più adatto a gestire un'elevata quantità di dati rispetto ai chip di front-end attualmente utilizzati in ATLAS.

4.7.2 Digital Chip Bottom

Il DCB contiene i circuiti che implementano le funzioni di controllo del chip e di processing dei dati. I principali blocchi che lo costituiscono sono rappresentati in figura 4.16, in cui è esplicitato il modo in cui esso comunica con l'ACB e con la matrice dei pixel.

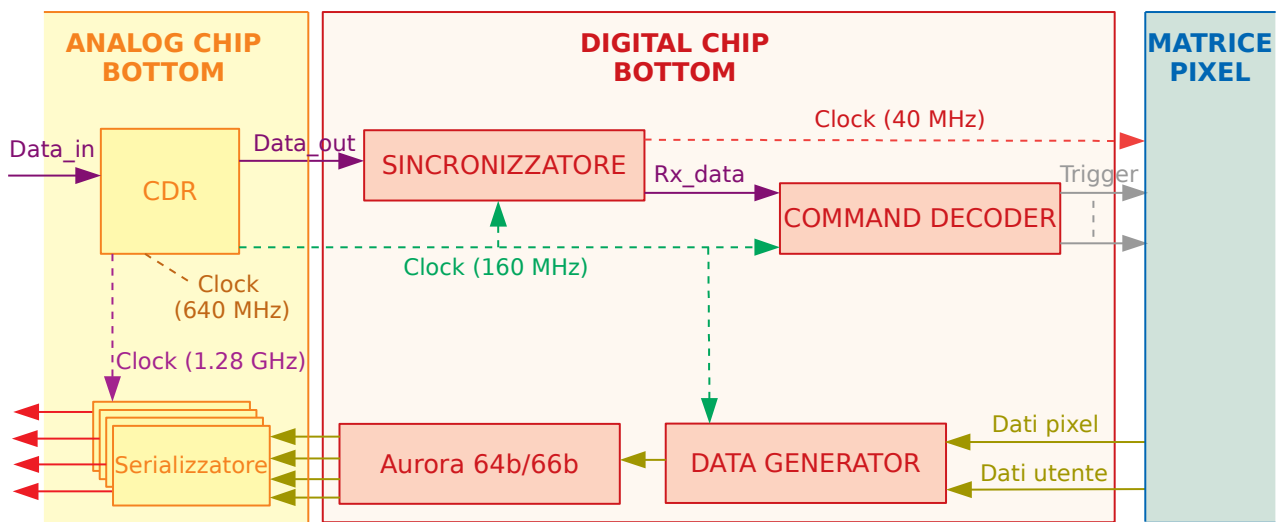


Figura 4.16: Illustrazione dei principali blocchi contenuti nel DCB e della sua connessione con l'ACB e con la matrice di pixel.

Il sincronizzatore genera il quarto clock utilizzato dal chip, a 40 MHz, a partire da quello a 160 MHz ricevuto dal CDR nell'ACB; tale clock è l'unico distribuito ai pixel. Il blocco indicato come *Command Decoder* decodifica i comandi di input, tra i quali quelli di trigger, di scrittura e lettura dei registri, e li trasmette alla matrice. I dati provenienti da ogni colonna di pixel ed i valori dei registri impostati dall'utente sono inviati, utilizzando il clock a 40 MHz, al *data generator*, il quale li organizza in pacchetti che trasmette, a 160 MHz, al blocco successivo. Tali dati sono codificati utilizzando il protocollo Aurora 64b/66b [44] ed inviati ai serializzatori contenuti nell'ACB attraverso 4 diverse linee. Questo trasferimento avviene con una frequenza di banda di 1.28 Gb/s, la quale può essere ridotta di un fattore 2, 4 oppure 8.

4.8 Risultati e sviluppi futuri

Il chip RD53A è stato utilizzato a partire dal 2017 per svolgere attività di test finalizzate a dimostrare la sua capacità di soddisfare i requisiti imposti dall'aumento di luminosità previsto durante la fase di HL–LHC. I risultati di tali analisi hanno guidato la realizzazione di un nuovo prototipo di chip, detto RD53B.

Al termine della fase di HL–LHC i due strati più interni dell'ITk avranno ricevuto una dose di radiazione pari a circa 12 MGy [45]; tale quantità è talmente elevata da non permettere l'utilizzo degli stessi chip durante l'intera fase di presa dati. L'obiettivo della collaborazione RD53 è quello di realizzare un chip che possa sostenere una dose di radiazione pari a 5 MGy; per questo motivo ITk è stato progettato in maniera tale che i due strati più interni possano essere sostituiti dopo aver raccolto una statistica corrispondente a 2000 fb^{-1} , pari alla metà di quella complessiva (come introdotto nella sezione 2.7). I test condotti su RD53A suggeriscono che RD53B potrebbe sostenere dosi di 10 MGy, raggiungendo dunque l'obiettivo prefissato.

Il chip RD53B è stato realizzato in due varianti per soddisfare le diverse richieste imposte dalle collaborazioni ATLAS e CMS; si distinguono così RD53B-ATLAS, anche detto ITK-v1, e RD53B-CMS, indicato anche come CROC-v1. Una prima differenza tra le due versioni è rappresentata dall'elettronica di front-end implementata: ATLAS ha deciso di utilizzare quella differenziale, CMS quella lineare. Inoltre, i due chip hanno diverse dimensioni: la matrice di pixel di RD53B-ATLAS è costituita da 400 colonne e 384 righe, mentre quella di RD53B-CMS ha 432 colonne e 336 righe. Entrambi i chip hanno dimensioni pari a quelle definitive, a differenza di RD53A.

La struttura di RD53B differisce da quella di RD53A, illustrata in figura 4.1, soltanto per il numero di righe di pixel. La logica di processing dei dati, schematizzata in figura 4.17, è analoga a quella del chip precedente. L'unica differenza rispetto a RD53A è rappresentata

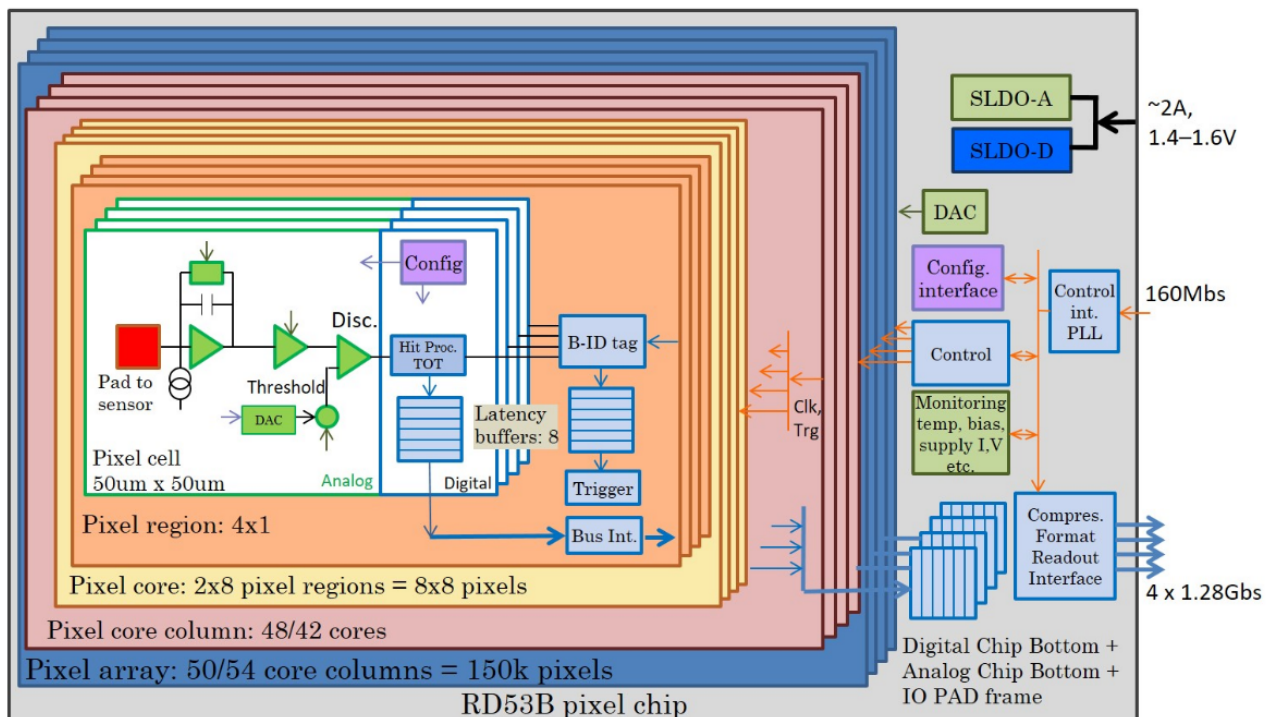


Figura 4.17: Schematizzazione della logica di processing dei dati in RD53B. Da [45].

dall'utilizzo della sola logica DBA per la lettura dei dati: 8 contatori di latenza da 9 bit ciascuno sono contenuti all'interno di una regione comune a 4 pixel, disposti su un'unica riga. Al termine del massimo tempo di latenza consentito (pari a quello di RD53A), i valori di ToT non nulli sono inviati al DCB nel caso in cui si sia ricevuto un segnale di trigger.

I numeri di colonne e righe di core indicati in figura sono relativi alla versione RD53B-ATLAS (50 e 48) e RD53B-CMS (54 e 42).

Capitolo 5

Procedura di calibrazione del chip RD53A

Prima di essere installato all'interno di un rivelatore, un generico chip di lettura è sottoposto a processi di test, caratterizzazione e calibrazione in laboratorio. Lo scopo di questa fase è quello di trovare il migliore insieme di parametri che definiscono il comportamento del chip. Ciò avviene eseguendo ciclicamente delle scansioni che consistono nell'iniettare nel front-end una quantità nota di carica o un impulso digitale che simuli la ricezione di un segnale sopra soglia, analizzare la risposta del chip ed eventualmente modificare i suoi parametri di configurazione in risposta ai dati ricevuti, in modo da ottenere il comportamento desiderato. Dopo aver calibrato il chip, il suo funzionamento può essere verificato in laboratorio sfruttando raggi cosmici o sorgenti radioattive.

In questo capitolo sono descritti il setup sperimentale utilizzato per la caratterizzazione del chip RD53A in esame, allestito nel Laboratorio INFN "Costruzione grandi apparati" presso il Dipartimento di Matematica e Fisica "E. De Giorgi", le caratteristiche principali del software adoperato e le procedure di configurazione del chip.

5.1 Setup sperimentale

Il setup sperimentale utilizzato nel presente lavoro di tesi per la caratterizzazione del chip RD53A è schematizzato in figura 5.1; il chip analizzato è contraddistinto dal numero seriale (*Serial Number*) 0x1995, e pertanto sarà indicato nel seguito come RD53A - SN: 0x1995.

Il chip analizzato è connesso ad un sensore di spessore attivo pari a 150 μm con pixel di dimensioni 25 $\mu\text{m} \times 100 \mu\text{m}$. La geometria utilizzata è quella 3D in configurazione 1E, realizzata a partire da un bulk di tipo *p*. Il sensore può essere alimentato tramite il generatore di tensione Keithley 6487, il quale funge anche da picoamperometro per la misura della corrente di leakage.

Come introdotto nella sezione 4.2 il chip RD53A è montato sulla Single Chip Card. Un generatore di tensione Agilent E3631A è collegato ad uno dei due connettori dedicati all'alimentazione del chip di cui è dotata la SCC; il secondo può essere sfruttato per alimentare serialmente diversi chip, come indicato in figura 4.4, ma non è stato utilizzato in questo lavoro.

La SCC dispone inoltre di due porte adoperate per la connessione dati, effettuata tramite cavi DisplayPort, ad una scheda detta *Ohio card*; essa è dotata di quattro porte, e ciò rende possibile il suo collegamento con più chip. Il compito della scheda Ohio è quello di interfacciare la SCC con la FPGA (*Field-Programmable Gate Array*) utilizzata per la lettura del chip. L'FPGA è montata su una scheda detta $\pi\text{LUP card}$ (*Pixel detector high Luminosity Lhc UPgrade*, [46]); questa comunica con il processore di un computer tramite il bus seriale PCIe (*Peripheral Component*

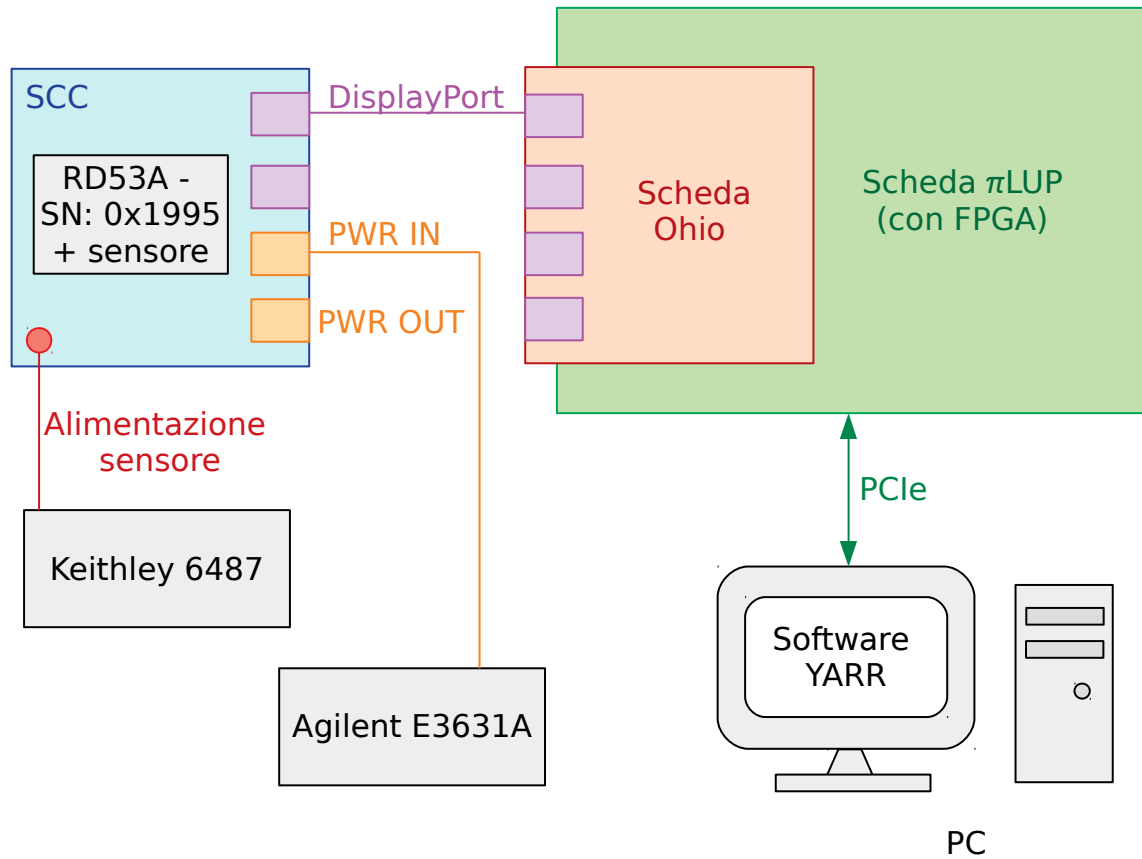


Figura 5.1: Schematizzazione del setup utilizzato durante il presente lavoro di tesi per la caratterizzazione del chip RD53A.

Interconnect express). Sul PC è installato il software YARR⁽¹⁾ (Yet another Rapid Readout, [47]), sul quale si basa il sistema di acquisizione dati (Data AcQuisition, DAQ) utilizzato in questo lavoro di tesi e descritto nella sezione 5.3.

5.2 Operazioni preliminari

Prima di poter adoperare il chip è necessario verificare che le tensioni e le correnti necessarie per il suo corretto funzionamento assumano i valori desiderati, stabiliti dai designer della collaborazione RD53. In figura 5.3 sono evidenziati in giallo i pin della SCC utilizzati per effettuare tale verifica; l'alimentazione fornita dal generatore Agilent E3631A è pari a 1.85 V.

Un primo valore da verificare è quello della corrente di riferimento utilizzata per gestire i DAC che forniscono l'alimentazione ai front-end; questa, indicata come I_{REF} , deve essere pari a 4 μ A. Per misurare tale corrente si collega il pin indicato in figura come I_{REF_IO} al generatore Keithley 6487, utilizzandolo come picoamperometro. I_{REF} è generata da un DAC a 4 bit, i quali possono essere selezionati fisicamente tramite i 4 pin indicati in figura come $I_{REF_trim\ bits}$. Il

⁽¹⁾Versione 1.1.0.

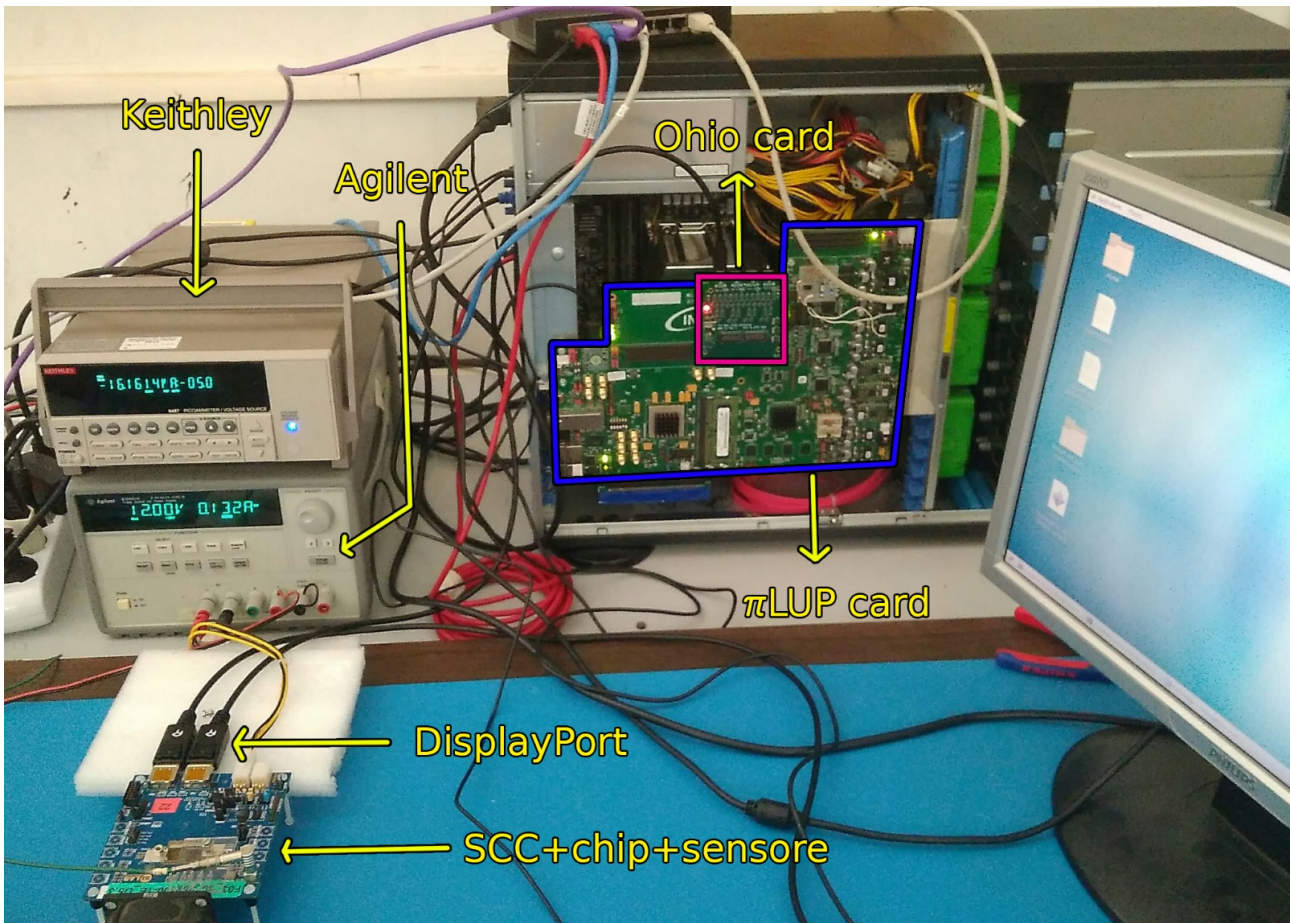


Figura 5.2: Foto del setup sperimentale.

valore di I_{REF} più vicino a quello ideale è ottenuto selezionando i connettori 1 e 2 ed è pari a $3.95 \mu A$.

Una seconda quantità da verificare è la tensione indicata come V_{REF_ADC} , essenziale per il corretto esito delle procedure di tuning delle soglie descritte nelle sezioni 5.5.3, 5.5.4 e 5.5.6. Questa è misurata utilizzando un multimetro Fluke 287 attraverso il pin indicato in figura come *Test V_{REF_ADC}* , considerando come riferimento la massa indicata come *GND*. Il valore ideale della tensione V_{REF_ADC} è pari a $0.9 V$; esso è regolato modificando il corrispondente parametro contenuto nel file di configurazione del chip.

Nel presente lavoro di tesi il chip RD53A è stato alimentato in modalità attiva, utilizzando i regolatori ShuLDO (si veda la sezione 4.3). Pertanto, è necessario verificare che essi forniscano in uscita le tensioni volute, pari a $1.2 V$; queste sono utilizzate per l'alimentazione delle componenti analogiche (V_{DDA}) e digitali (V_{DDD}) dei front-end. Le due tensioni sono misurate con un multimetro tramite i pin indicati come *Test point V_{DDA}/V_{DDD}* ; qualora siano molto distanti dal valore desiderato possono essere modificate agendo sui relativi parametri contenuti nel file di configurazione. I valori ottenuti al termine di questa procedura sono $V_{DDA} = (1.2037 \pm 0.0001) V$ e $V_{DDD} = (1.2029 \pm 0.0001) V$.

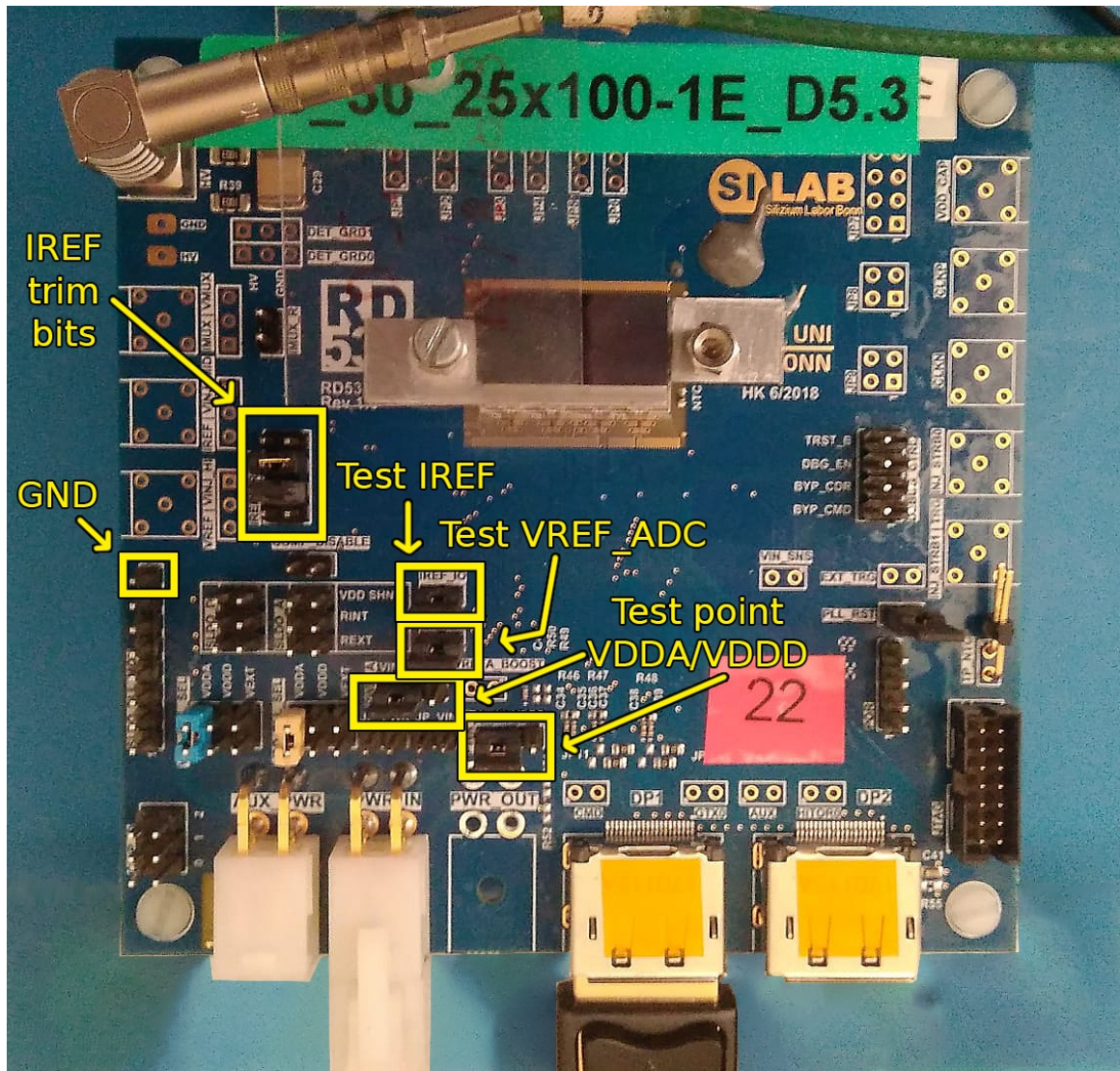


Figura 5.3: Indicazione dei pin della SCC utilizzati per verificare il valore di tensioni e correnti necessarie per il corretto funzionamento del chip (in giallo).

5.3 Il sistema di lettura YARR

Il sistema di lettura YARR è stato sviluppato per la calibrazione di chip di front-end. Esso nasce in risposta all'esigenza di modificare il sistema di acquisizione dati utilizzato per la lettura dei chip FEI4 installati nell'IBL, ma è stato progettato in modo tale da essere adattato ad altri chip. Questa caratteristica ha reso possibile il suo utilizzo nel presente lavoro di tesi, centrato intorno al chip RD53A. La differenza tra la logica utilizzata per la lettura di FEI4 e quella alla base del sistema YARR è rappresentata in figura 5.4. Ogni modulo del rivelatore a pixel dell'IBL è collegato ad una scheda detta *Read Out Driver* (ROD). Essa contiene una FPGA che si occupa di eseguire le scansioni durante la fase di calibrazione del chip (*Scan engine*) ed effettua una prima manipolazione dei dati. La comunicazione tra la FPGA ed il processore del PC, infatti, è troppo lenta affinché essi siano trasferiti dall'una all'altro; per questo motivo i dati sono organizzati in istogrammi all'interno della FPGA (*Histogrammer*) e successivamente inviati alla CPU (*Central Processing Unit*) del computer, perdendo così parte dell'informazione.

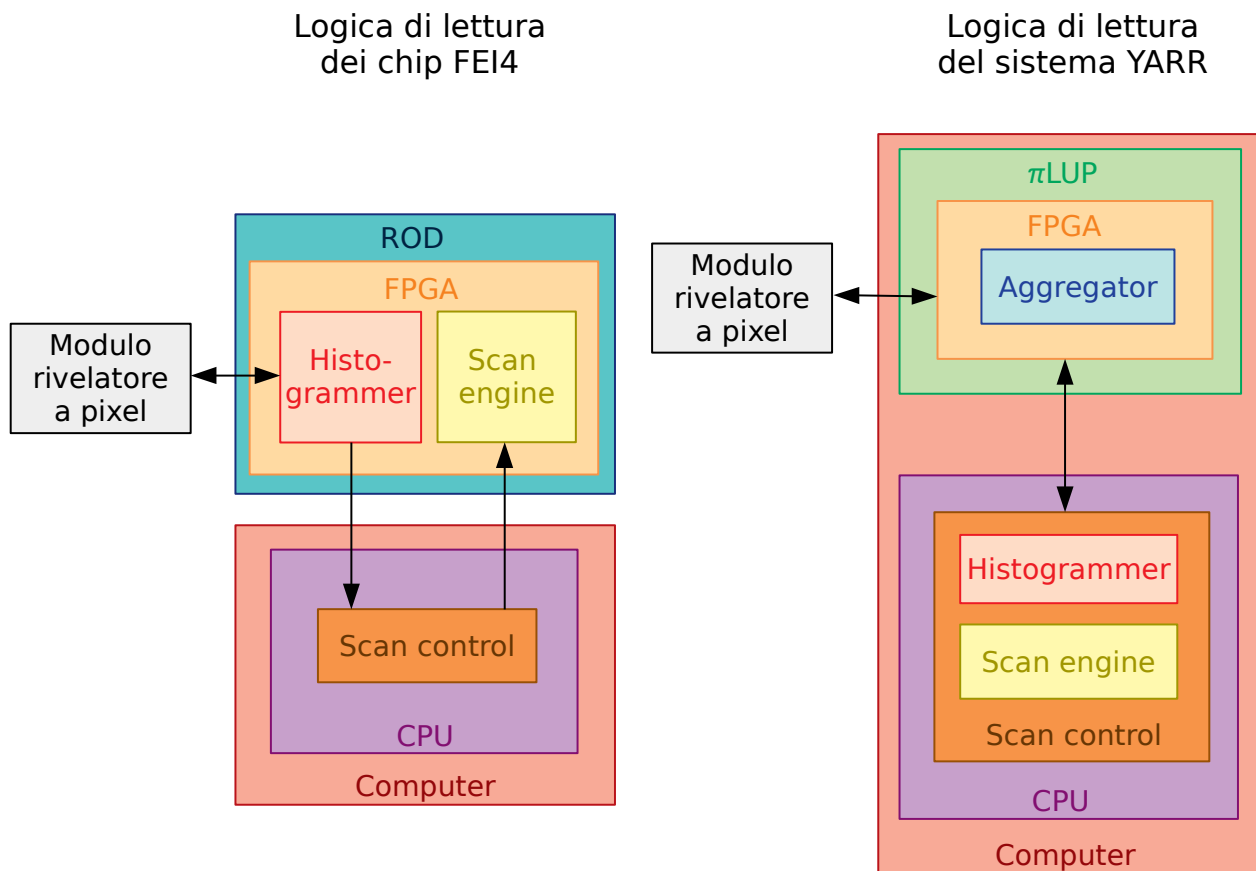


Figura 5.4: Differenza tra la logica utilizzata per la lettura dei chip di front-end FEI4 installati nell'IBL (sinistra) e quella alla base del sistema YARR (destra).

Il processore controlla le scansioni eseguite dalla FPGA (*Scan control*) ed analizza gli istogrammi ricevuti.

L'idea alla base del progetto YARR consiste nel disaccoppiare il software dal firmware della FPGA, facendo svolgere al primo la maggior parte delle funzioni. Poiché il firmware è estremamente interconnesso con l'hardware utilizzato, il suo disaccoppiamento dal software è il motivo per il quale quest'ultimo può essere adoperato per la lettura di diversi chip. Ciò consente inoltre di ottimizzare il firmware ed il software in maniera indipendente, impiegando dunque meno tempo, e di semplificare la transizione dall'hardware utilizzato in laboratorio per la calibrazione dei chip a quello del rivelatore, inevitabilmente più complesso. Infine il software sul quale si basa il sistema YARR, scritto in C++, è più facilmente accessibile rispetto al firmware della FPGA, e ciò rende preferibile il suo utilizzo.

Nel sistema YARR l'FPGA svolge dunque un minor numero di funzioni. Essa, direttamente collegata al computer e non più contenuta in schede connesse ai singoli moduli del rivelatore, funge da semplice interfaccia di input/output, aggregando i dati letti dal front-end (*Aggregator*) ed inviandoli al processore. Quest'ultimo si occupa della gestione e dell'effettiva esecuzione delle scansioni necessarie alla calibrazione del chip ed anche dell'organizzazione dei dati ricevuti

in istogrammi. In tal modo nella CPU sono contenute tutte le informazioni provenienti dall'elettronica di front-end e ciò rende possibile lo svolgimento di un'analisi più dettagliata. Il sistema YARR, quindi, trae vantaggi dal notevole progresso della tecnologia che consente il trasferimento di una maggiore quantità di dati tra l'FPGA e la CPU mediante il bus seriale PCIe, sviluppato nel corso degli ultimi anni per gestire le schede video di ultima generazione.

5.4 Logica delle scansioni

Come accennato nella sezione 5.3, la calibrazione del chip avviene eseguendo una serie di scansioni. Ciascuna di esse è costituita da diversi processi, indicati in figura 5.5.

Lo Scan engine legge il file contenente i parametri di configurazione del chip e si occupa dell'esecuzione di una successione di cicli (descritti con maggiore dettaglio nella sezione 5.4.1), all'interno di ciascuno dei quali è possibile individuare quattro funzioni principali:

- *Init()*, eseguita un'unica volta all'inizio del ciclo;
- *ExecPart1()*, rappresenta la prima parte del ciclo e può essere eseguita più volte;
- *ExecPart2()*, costituisce la seconda parte del ciclo ed è anch'essa eseguibile più volte;
- *End()*, eseguita un'unica volta al termine del ciclo.

La modalità di esecuzione dei cicli è schematizzata in figura 5.6. Nel caso in cui la variabile booleana *has inner* ha valore 1 la funzione *ExecPart1()* richiama l'esecuzione di un nuovo ciclo. In base al valore assunto dalla seconda variabile booleana *done*, al termine dell'esecuzione della funzione *ExecPart2()* si ripete quella di *ExecPart1()* oppure è eseguita la funzione *End()*, terminando così il ciclo. Tale struttura risulta essere facilmente modificabile dall'utente, il quale può variare l'ordine ed il numero di cicli da eseguire all'interno di ogni scansione. Ogni ciclo eseguito dallo Scan engine produce un frammento di dati indicato come *raw data*, il quale è salvato in un'area di memoria detta *Bookkeeper*.

Il livello di processamento successivo è costituito da una serie di sottoprocessi (*thread*) indicati come *Data processor*, i quali lavorano in parallelo; questi prelevano i raw data, li decodificano e costruiscono degli eventi, in numero pari a quello dei chip di front-end collegati alla scheda π LUP. Ciascuno di essi è ulteriormente processato da due thread successivi; il primo, detto *Histogrammer*, preleva gli eventi e li organizza in istogrammi, i quali sono poi analizzati dal secondo thread, detto *Analyser*.

Alcune delle scansioni eseguite durante la fase di caratterizzazione del chip sono ripetute al variare dei valori assunti da alcuni registri di configurazione, con l'obiettivo di individuare quelli che permettono di ottenere il comportamento desiderato. Ciò comporta la necessità di inviare dei feedback dal thread di *Analyser* a quello di *Scan engine*, i quali sono dunque gli unici due tra i quali è stabilita una comunicazione. La scansione termina nel momento in cui è individuata la configurazione ottimale del chip, la quale viene salvata in memoria.

5.4.1 Cicli eseguiti nelle scansioni

I cicli che possono essere eseguiti durante la procedura di configurazione del chip RD53A sono in totale 6; di seguito sono fornite le informazioni relative ad ognuno di essi.

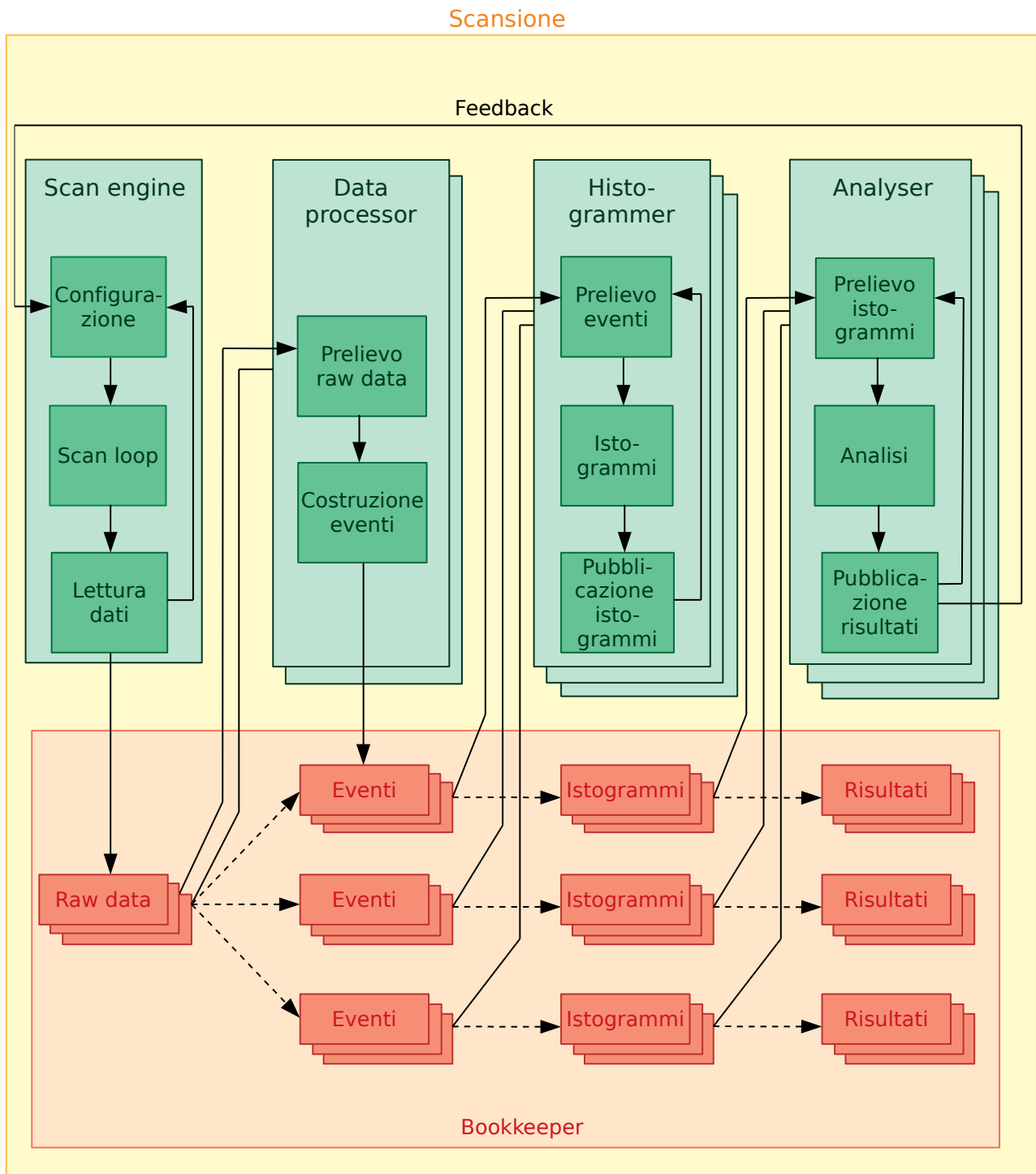


Figura 5.5: Schematizzazione della logica delle scansioni eseguite dal software YARR.

Rd53aTriggerLoop

Definisce le modalità di iniezione e di generazione del segnale di trigger.

Al suo interno è possibile specificare il numero di iniezioni da effettuare, la frequenza di invio del segnale di trigger, la distanza temporale tra l'iniezione e la generazione del trigger (espressa

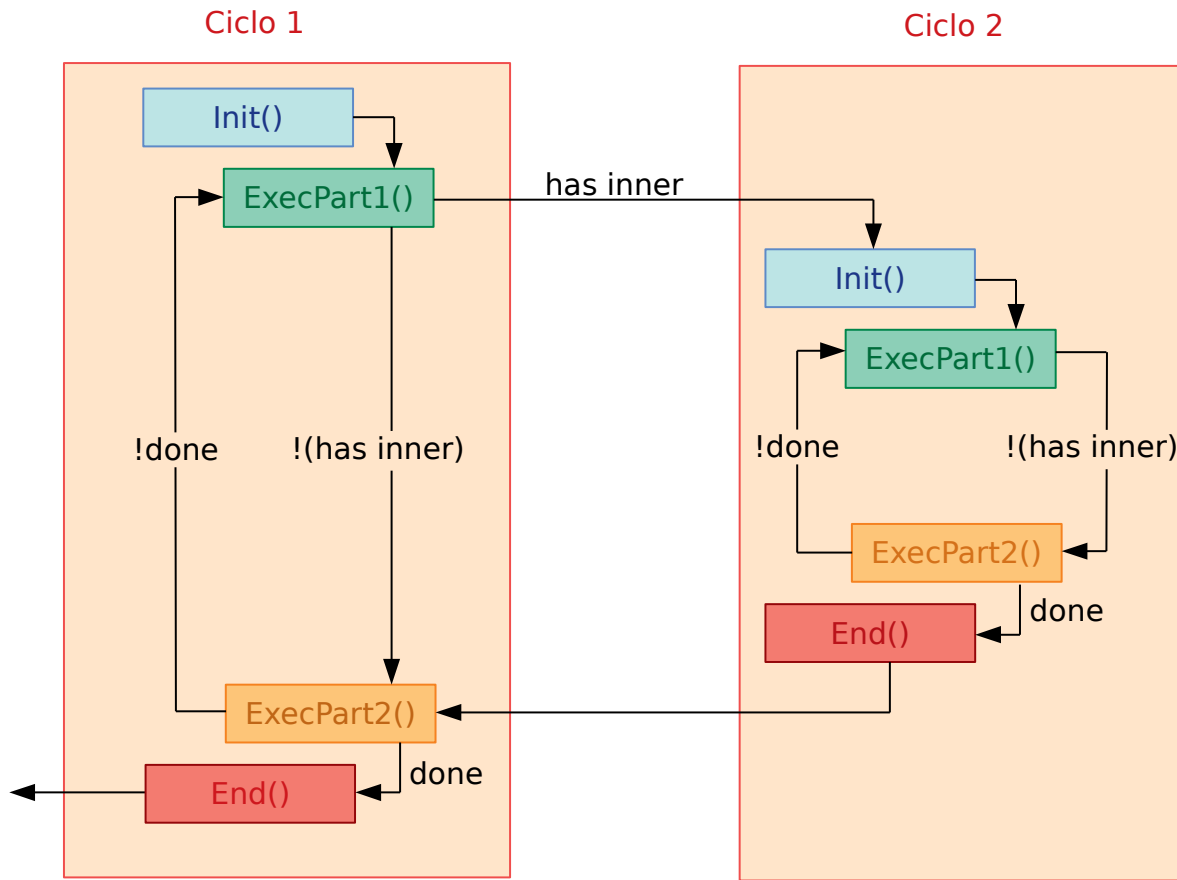


Figura 5.6: Logica di esecuzione dei cicli da parte di una generica scansione eseguita utilizzando il software YARR durante la procedura di configurazione del chip.

in unità di bc), la durata del ciclo, l'abilitazione della ricezione di un trigger esterno e dei circuiti di iniezione.

Rd53aMaskLoop

Seleziona determinati pixel da abilitare per l'iniezione e la configurazione; l'applicazione di una maschera (ovvero l'attivazione di una porzione dell'intera matrice di pixel) è necessaria per evitare picchi nella corrente di assorbimento che altererebbero la risposta del chip. Per poter individuare i pixel questi sono numerati all'interno di ogni colonna di core secondo la relazione

$$n_pixel = (riga_core \cdot 64) + [(colonna_pixel + riga_core \% 8) \% 8] \cdot 8 + riga_pixel \% 8$$

in cui n_pixel indica la posizione del pixel all'interno della colonna di core, $riga_core$ la posizione del core al quale appartiene il pixel considerato all'interno della matrice, $colonna_pixel$ e $riga_pixel$ quella del pixel nella matrice. Le righe dei core e dei singoli pixel all'interno della matrice sono numerate in modo crescente dall'alto verso il basso, le colonne da sinistra verso destra. La numerazione descritta è illustrata graficamente in figura 5.7, nella quale è rappre-

sentato il primo core della matrice, per il quale $\text{riga_core} = \text{colonna_core} = 0$; in rosso sono indicati i valori di riga_pixel e colonna_pixel .

	0	1	2	3	4	5	6	7
0	0	8	16	24	32	40	48	56
1	1	9	17	25	33	41	49	57
2	2	10	18	26	34	42	50	58
3	3	11	19	27	35	43	51	59
4	4	12	20	28	36	44	52	60
5	5	13	21	29	37	45	53	61
6	6	14	22	30	38	46	54	62
7	7	15	23	31	39	47	55	63

Prima maschera

Seconda maschera

Terza maschera

Quarta maschera

Quinta maschera

Sesta maschera

Settima maschera

Ottava maschera

Figura 5.7: Numerazione dei pixel nel primo core della matrice ed indicazione delle maschere applicate nel caso in cui $\text{min} = 0$, $\text{max} = 8$, $\text{step} = 8$.

All'interno del ciclo è possibile specificare il numero identificativo del primo pixel da abilitare (min), la distanza tra i pixel da abilitare nella stessa iterazione (max) ed il numero di maschere che è necessario applicare affinché al termine dell'ultimo ciclo l'intero core sia stato abilitato (step). Si consideri, a titolo esemplificativo, il caso in cui $\text{min} = 0$, $\text{max} = 8$. In tal caso, per abilitare l'intero core è necessario applicare iterativamente 8 maschere ($\text{step} = 8$), indicate in figura 5.7.

Rd53aCoreColLoop

Seleziona le colonne all'interno delle quali abilitare i pixel indicati nel ciclo Rd53aMaskLoop.

All'interno del ciclo è possibile specificare gli estremi dell'intervallo al quale appartengono le colonne di core da attivare (min e max), la distanza tra quelle da abilitare nello stesso ciclo (step) ed il numero di maschere da applicare (nSteps). Si consideri, a titolo esemplificativo, il caso in cui $\text{min} = 0$, $\text{max} = 50$, $\text{step} = 10$, $\text{nSteps} = 10$. In questo caso, la prima maschera applicata consentirà l'abilitazione delle colonne di core numerate come 0, 10, 20, 30, 40, 50; la maschera successiva selezionerà i core nelle colonne 1, 11, 21, 31, 41 e così via. Affinché l'intera matrice di pixel sia stata abilitata al termine dell'esecuzione del ciclo è necessario applicare consecutivamente 10 maschere; da ciò deriva il valore assunto dalla variabile nSteps .

Rd53aGlobalFeedback

Questo ciclo è utilizzato per modificare il valore di un determinato registro globale sulla base del feedback ricevuto dall'algoritmo di analisi, come indicato in figura 5.5.

Al suo interno è possibile specificare il nome del registro sul quale agire, gli estremi del suo intervallo di appartenenza ed il passo della variazione.

Rd53aPixelFeedback

Ha una funzione analoga a quella svolta dal ciclo Rd53aGlobalFeedback, ma agisce sui valori dei registri locali.

Rd53aParameterLoop

I cicli precedentemente descritti possono essere eseguiti al variare del valore assunto da uno dei parametri di configurazione (quale, ad esempio, quello che definisce la quantità di carica iniettata all'ingresso del front-end). Il nome del parametro in questione, gli estremi del suo intervallo di appartenenza ed il passo della variazione sono indicati nel ciclo Rd53aParameterLoop.

5.5 Procedura di configurazione del chip

Le scansioni eseguite per configurare il chip nel modo desiderato possono essere distinte in due categorie: procedure di *tuning* o di *scan*. Nel primo caso i registri che definiscono i parametri caratteristici del chip sono modificati all'inizio di ogni iterazione sulla base del feedback ricevuto dall'algoritmo di analisi, come indicato in figura 5.5, ed il loro valore ottimale è salvato in memoria al termine della scansione; uno scan è invece utilizzato per misurare il valore effettivo di un parametro, senza alterare la configurazione del chip.

La procedura seguita per una corretta configurazione del chip è costituita dalle seguenti scansioni:

- scan digitale dell'intera matrice;
- scan analogico dell'intera matrice;
- tuning della soglia globale del front-end differenziale;
- tuning della soglia locale del front-end differenziale;
- tuning del CSA del front-end differenziale;
- tuning della soglia locale del front-end differenziale;
- tuning fine della soglia locale del front-end differenziale;
- tuning della soglia globale del front-end lineare;
- tuning della soglia locale del front-end lineare;
- tuning del CSA del front-end lineare;
- tuning della soglia locale del front-end lineare;
- tuning fine della soglia locale del front-end lineare;

- tuning della soglia (globale) del front-end sincrono;
- tuning del CSA del front-end sincrono;
- tuning della soglia (globale) del front-end sincrono;
- scan della soglia dell'intera matrice;
- scan del ToT dell'intera matrice.

Per evitare che un malfunzionamento di uno dei front-end influenzi anche gli altri due, le procedure di tuning sono eseguite separatamente per ognuno di essi, definendo opportunamente i valori dei parametri min e max nel ciclo di abilitazione delle colonne di core. La distinzione tra il tuning ed il tuning fine della soglia locale è dettata dall'algoritmo utilizzato per impostare la soglia locale, come è descritto nella sezione 5.5.6. Si può osservare che per il front-end sincrono non è eseguito alcun tuning della soglia locale, dal momento che la dispersione è ridotta utilizzando il meccanismo di auto-azzeramento descritto nella sezione 4.5.3.

Nelle prossime sezioni sono fornite maggiori informazioni in merito ad ognuna delle scansioni elencate. Inoltre, saranno illustrati i risultati di ciascuna procedura ottenuti sul chip RD53A - SN: 0x1995 letto nel Laboratorio INFN "Costruzione grandi apparati" mediante il setup descritto nella sezione 5.1.

5.5.1 Scan digitale

Per verificare il corretto funzionamento della componente digitale del front-end è possibile iniettare all'uscita del discriminatore degli impulsi che simulino la ricezione di un segnale sopra soglia, come introdotto nella sezione 4.4.

I cicli eseguiti durante questo scan sono Rd53aMaskLoop, Rd53aCoreColLoop e Rd53aTriggerLoop. Per disaccoppiare la componente digitale sotto esame da quella analogica, la soglia dei discriminatori nei tre diversi front-end è impostata ad un valore sufficientemente elevato da evitare che si possa generare un hit come conseguenza del rumore e non dell'iniezione digitale.

Lo scan digitale produce due mappe bidimensionali, nelle quali ogni bin corrisponde ad uno dei pixel del chip; un esempio è illustrato in figura 5.8, relativa ad uno scan in cui si sono effettuate 100 iniezioni. Nella mappa a sinistra, indicata come *Mappa degli hit*, è indicato il numero di hit registrati da ogni pixel; è possibile osservare come nel caso illustrato ciascun pixel del chip RD53A - SN: 0x1995 abbia rivelato tutte le iniezioni effettuate. Nella mappa a destra, detta *Mappa dei pixel funzionanti*, ad ogni pixel è associato il valore 1 nel caso in cui il numero di hit registrati sia pari a quello delle iniezioni effettuate, 0 altrimenti. Tale mappa può eventualmente essere utilizzata per disabilitare negli scan successivi i pixel che hanno risposto in modo non adeguato; per questo motivo è spesso indicata come *enable map*. Poiché nello scan eseguito tutti i pixel hanno manifestato un comportamento corretto, a ciascuno di essi è associato il valore 1.

5.5.2 Scan analogico

Utilizzando il circuito di iniezione posto all'ingresso dell'amplificatore di ciascun front-end, illustrato in figura 4.4, è possibile iniettare una certa quantità di carica e contare il numero di hit registrati da ogni pixel del chip RD53A - SN: 0x1995.

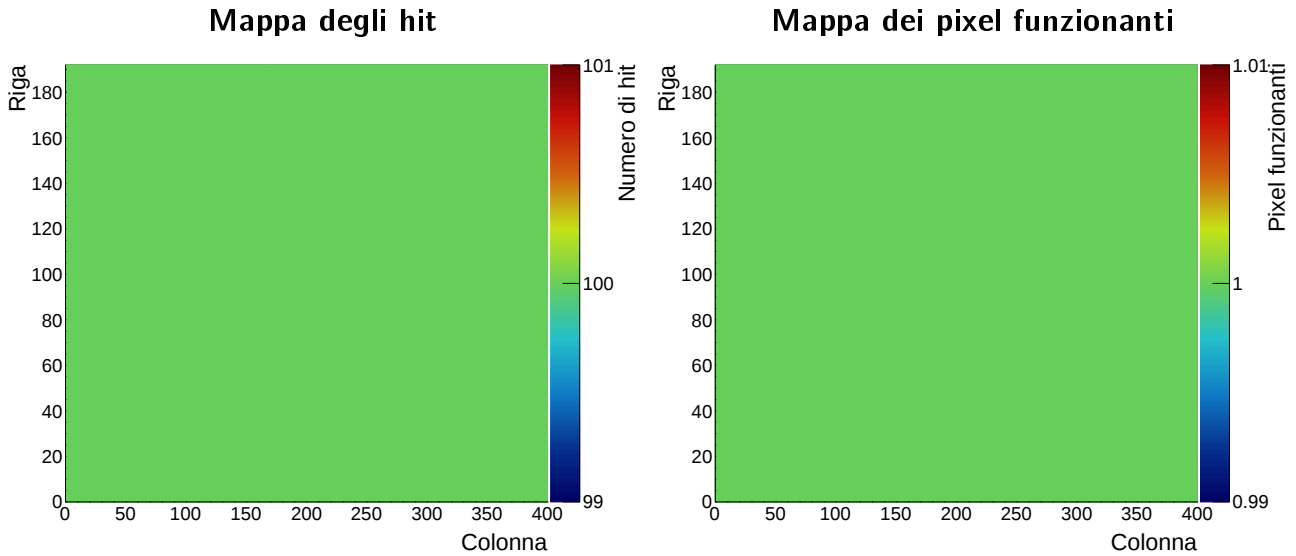


Figura 5.8: Esempio di mappe bidimensionali del chip RD53A - SN: 0x1995 prodotte al termine di uno scan digitale in cui si sono effettuate 100 iniezioni.

I cicli eseguiti durante questa scansione sono gli stessi utilizzati per lo scan digitale. La carica iniettata è sufficientemente elevata affinché sia sicuramente superiore alla soglia, determinata dal valore del corrispondente registro di configurazione salvato in memoria. Questa scansione produce delle mappe analoghe a quelle fornite dallo scan digitale, rappresentate in figura 5.9. È

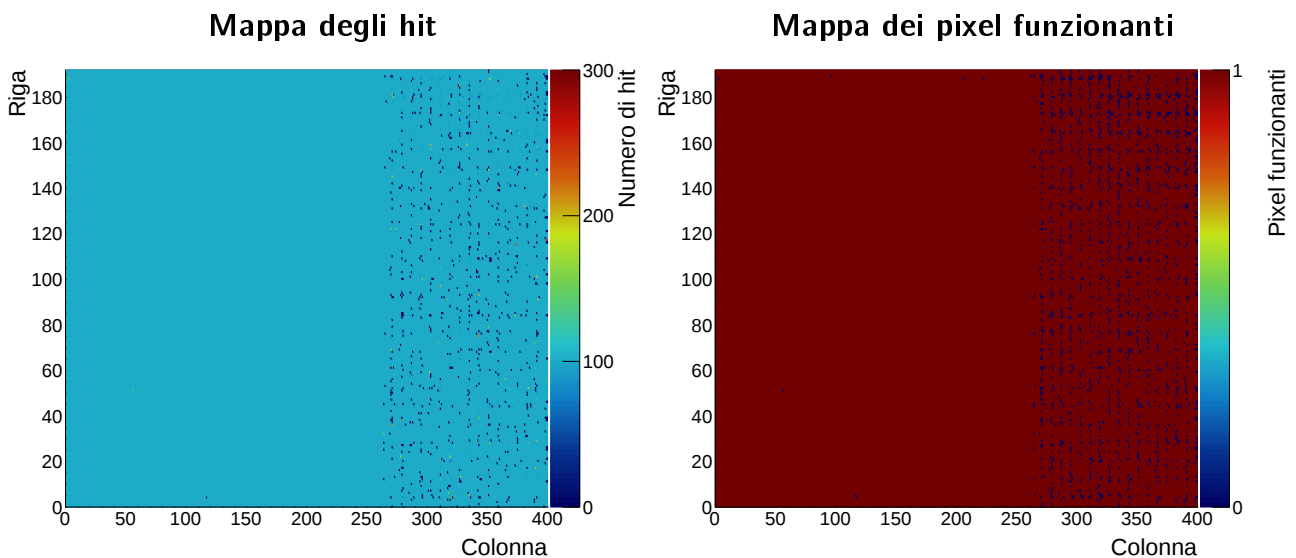


Figura 5.9: Esempio di mappe bidimensionali del chip RD53A - SN: 0x1995 prodotte al termine di uno scan analogico in cui si sono effettuate 100 iniezioni.

possibile osservare come in questo caso molti pixel appartenenti alla sezione differenziale della matrice manifestano un comportamento scorretto; ad essi è pertanto associato il valore 0 nella enable mask (a destra in figura 5.9).

Combinando i risultati dello scan analogico con quelli dello scan digitale è possibile individuare i pixel per i quali la componente analogica del front-end manifesta un comportamento scorretto; infatti, un suo malfunzionamento influenza la risposta della componente digitale, ma non viceversa. Nel caso illustrato tutti i pixel hanno risposto correttamente all'iniezione digitale e pertanto

il problema osservato nel front-end differenziale deve essere attribuito alla sua componente analogica. Tale problema è noto, ed è stato identificato successivamente alla sottomissione del chip come un difetto di design che causa l'introduzione di capacità parassite in un sottoinsieme di pixel all'interno di ogni core. La disposizione di tali pixel è la stessa per ogni core ed è illustrata in rosso in figura 5.10.

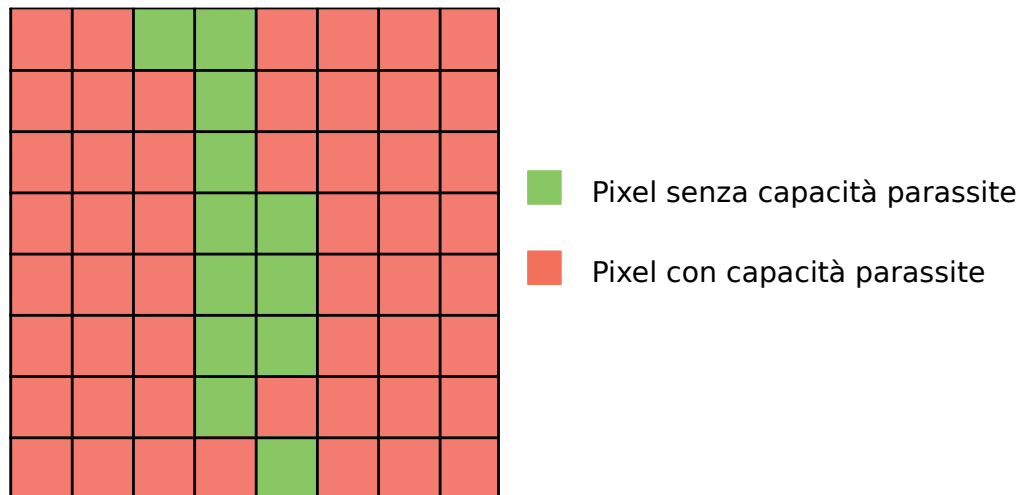


Figura 5.10: Rappresentazione dei pixel nei quali, a causa di un errore di progettazione, si manifestano correnti parassite all'interno di un core del front-end differenziale del chip RD53A.

Le correnti parassite distorcono il segnale rivelato all'interno di un pixel, rendendo scorretta la sua risposta. Per risolvere tale problema è possibile escludere dall'analisi i pixel indicati in rosso in figura 5.10 applicando una maschera a ciascun core; questa sarà indicata nel seguito come *maschera differenziale*.

Per verificare se i pixel che rispondono in modo scorretto siano quelli nei quali si individuano capacità parassite si realizza una nuova mappa dei pixel funzionanti, assegnando valore 0 ai pixel abilitati dalla maschera differenziale che registrano un numero di hit diverso da quello atteso, valore 1 ai pixel abilitati che registrano 100 hit e valore 2 ai pixel disabilitati. Il risultato ottenuto è illustrato in figura 5.11. Si osserva come all'interno del front-end differenziale non ci siano pixel ai quali è associato il valore 0; ciò vuol dire che tutti quelli che hanno manifestato un comportamento scorretto nella figura 5.9 sono quelli all'interno dei quali si osservano capacità parassite.

5.5.3 Tuning della soglia globale

L'obiettivo di questa scansione è quello di impostare la soglia globale dei discriminatori al valore desiderato.

In aggiunta ai cicli Rd53aMaskLoop, Rd53aCoreColLoop e Rd53aTriggerLoop, il tuning della soglia è ottenuto eseguendo anche Rd53aGlobalFeedback; ciò che viene modificato in ogni iterazione è il registro contenente i bit convertiti dai DAC nella tensione di soglia del discriminatore. In corrispondenza di ciascuno dei valori assunto da tale registro è iniettata una quantità di carica pari alla soglia desiderata ed è misurato il numero di hit registrato da ogni pixel. Le

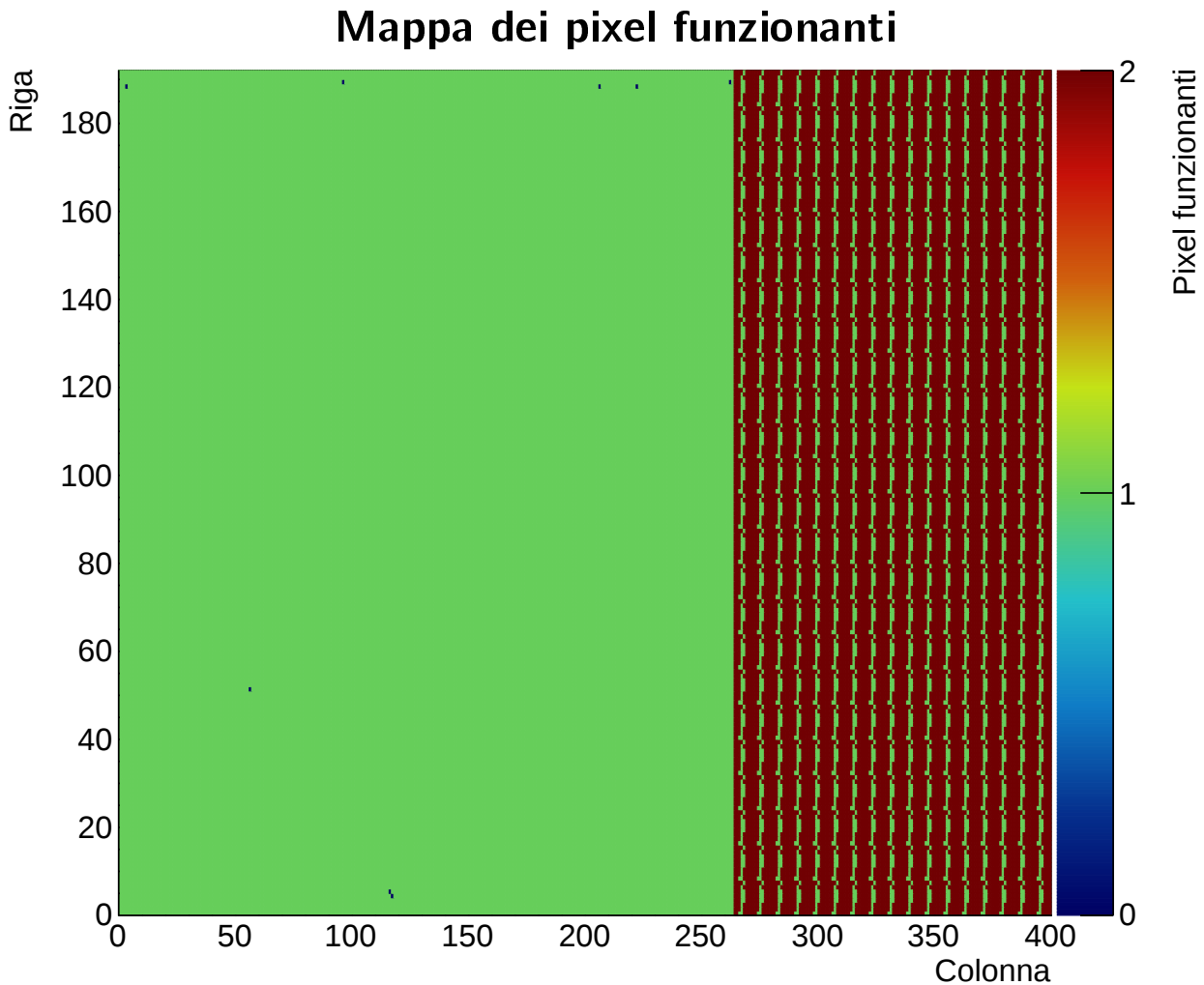


Figura 5.11: Mappa dei pixel funzionanti del chip RD53A - SN: 0x1995 relativa ad uno scan analogico nella quale si è associato valore 0 ai pixel abilitati dalla maschera differenziale che registrano un numero di hit diverso da quello atteso, valore 1 ai pixel abilitati che registrano 100 hit e valore 2 ai pixel disabilitati.

iterazioni sono definite in modo tale da diminuire gradualmente la soglia, iniziando da un valore presumibilmente superiore rispetto a quello della carica iniettata; di conseguenza, il numero di hit registrato da ogni pixel è inizialmente nullo. Assumendo che la distribuzione delle soglie di ogni pixel abbia un andamento gaussiano intorno al valore medio, il numero di hit osservati è pari al 50% di quello atteso in corrispondenza di un valore di soglia coincidente con quello della carica iniettata, e cioè della soglia desiderata. La procedura di tuning termina nel momento in cui si verifica tale condizione ed i valori dei bit corrispondenti sono salvati nel file di configurazione.

5.5.4 Tuning della soglia locale

La dispersione della soglia nel front-end differenziale ed in quello lineare è ridotta eseguendo un tuning della soglia locale.

L'unica differenza rispetto alla scansione precedente è rappresentata dal fatto che l'algoritmo Rd53aGlobalFeedback è sostituito da quello Rd53aPixelFeedback, il quale agisce sui registri

contenenti i bit dei DAC che forniscono la corrente (nel caso del front-end lineare) e la tensione (per il differenziale) necessarie per ridurre la dispersione di soglia.

5.5.5 Tuning del CSA

L'obiettivo del tuning del CSA è quello di modificare la relazione di proporzionalità tra la carica all'ingresso dell'amplificatore di ogni front-end ed il ToT ad essa corrispondente.

I cicli eseguiti durante questa scansione sono Rd53aMaskLoop, Rd53aCoreColLoop, Rd53aTriggerLoop e Rd53aGlobalFeedback. Quest'ultimo agisce sul parametro che determina la scarica del condensatore integrato nel CSA: la corrente indicata come $I_{ff}^{(2)}$ in figura 4.11 per il front-end differenziale, quella di Krummenacher per gli altri due, indicata come I_K nelle figure 4.10 e 4.13⁽³⁾. Al termine di ogni iterazione sui parametri che definiscono tali correnti è calcolata la media dei valori di ToT registrati da ciascun pixel, effettuata sul numero di iniezioni; si ricorda che il ToT è calcolato in unità di bc, con $1 \text{ bc} = 25 \text{ ns}$. Tale informazione è sfruttata per realizzare le distribuzioni delle medie relative ad ogni front-end. La procedura di tuning termina nel momento in cui le medie di queste distribuzioni coincidono con il valore di ToT voluto, specificato nel momento in cui si avvia la scansione.

In figura 5.12 sono rappresentate le distribuzioni ottenute al termine di una procedura di tuning nella quale il ToT desiderato per una carica di 10000 e^- è pari a 8 bc, ovvero $200 \text{ ns}^{(4)}$: il grafico in alto a sinistra è relativo al front-end (FE) sincrono, quello in alto a destra al lineare e quello in basso al differenziale. In tutti i casi il valor medio è vicino, entro la dispersione della distribuzione, a quello richiesto. Il minor numero di eventi nell'ultimo istogramma è dovuto all'applicazione della maschera differenziale.

Poiché la calibrazione della corrente di scarica del condensatore modifica la forma del segnale, la procedura appena descritta è seguita da un nuovo tuning della soglia (locale per i front-end differenziale e lineare, globale per quello sincrono).

5.5.6 Tuning fine della soglia locale

La scansione indicata come tuning fine ha l'obiettivo di ridurre ulteriormente la dispersione di soglia utilizzando un algoritmo diverso rispetto a quello descritto nella sezione 5.5.4.

I cicli eseguiti durante questa procedura sono Rd53aMaskLoop, Rd53aCoreColLoop, Rd53aTriggerLoop, Rd53aPixelFeedback e Rd53aParameterLoop. Rd53aPixelFeedback agisce, come nel tuning descritto nella sezione 5.5.4, sui valori dei bit di aggiustamento della soglia locale; l'ultimo ciclo, invece, modifica la carica iniettata⁽⁵⁾. Per ognuno dei valori assunti dai bit di aggiustamento di un pixel viene realizzata una curva S, la quale rappresenta la percentuale di hit registrati rispetto al numero di iniezioni effettuate in funzione della carica iniettata. In figura 5.13 è rappresentato un esempio di curva S, nella quale la carica è espressa in unità di carica elementare. Nel caso ideale, le curva assumerebbe valore 0% per valori di carica inferiori rispetto

⁽²⁾ *DiffVff* nel software YARR.

⁽³⁾ Rispettivamente *LinKrumCurr* e *SynclbiasKrum* nel software YARR.

⁽⁴⁾ Si sceglie di valutare le prestazioni del chip in corrispondenza di 10000 e^- perché questa quantità di carica corrisponde al più probabile numero di elettroni eccitati in banda di conduzione al passaggio di una particella carica al minimo di ionizzazione in $150 \mu\text{m}$ di Silicio, ovvero lo spessore dei sensori che saranno utilizzati nell'ITk. Per maggiori dettagli si veda la figura 6.15 nella sezione 6.8.

⁽⁵⁾ *InjVcalDiff* nel software YARR.

Distribuzioni ToT medio

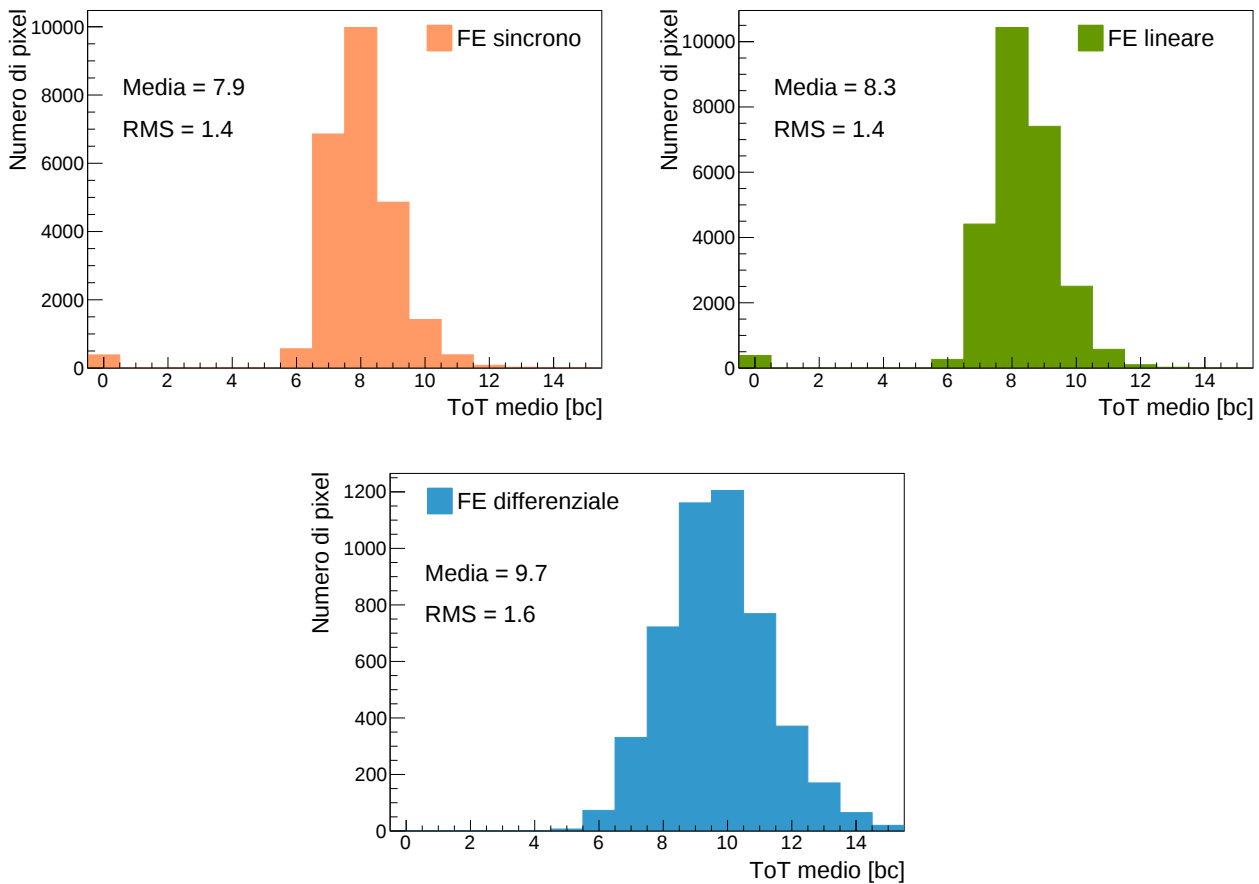


Figura 5.12: Risultato della procedura di tuning della corrente di scarica del CSA per il front-end sincrono (in alto a sinistra), lineare (in alto a destra) e differenziale (in basso) del chip RD53A - SN: 0x1995.

alla soglia e 100% viceversa. Il suo reale andamento, tuttavia, è distorto a causa del rumore elettronico che determina l'esistenza di un intervallo di valori di carica nel quale la percentuale di hit registrati è compresa tra 0% e 100%.

Dal momento che la fluttuazione del segnale del singolo pixel è gaussiana, è possibile eseguire un fit della curva S relativa a ciascun bit utilizzando una funzione degli errori; esso assumerà un valore pari al 50% in corrispondenza di una carica iniettata pari alla soglia desiderata, come illustrato in figura 5.13. Tale procedura viene realizzata per ognuna delle curve ottenute variando i valori dei bit di aggiustamento; queste sono 16 per ciascuno dei pixel nei quali è implementato il front-end lineare, 32 per ogni pixel nel front-end differenziale (si ricorda che il DAC utilizzato per la regolazione della soglia nella sezione lineare è a 4 bit, mentre nella sezione differenziale è presente un quinto bit che seleziona la soglia sulla quale agire tra le due presenti). Confrontando ciascuna soglia ottenuta dai fit delle curve S con la soglia desiderata, dichiarata nel momento dell'avvio della procedura di tuning, si individuano i valori ideali dei registri che regolano i DAC locali, i quali sono salvati nel file di configurazione.

Dalla sovrapposizione delle curve S associate ai valori ideali dei DAC di ciascun pixel è possibile ottenere una curva che descriva il comportamento complessivo dell'area del chip sottoposta al test; in figura 5.14 è riportato il risultato relativo all'esecuzione di 50 iniezioni di carica nel

Curva S

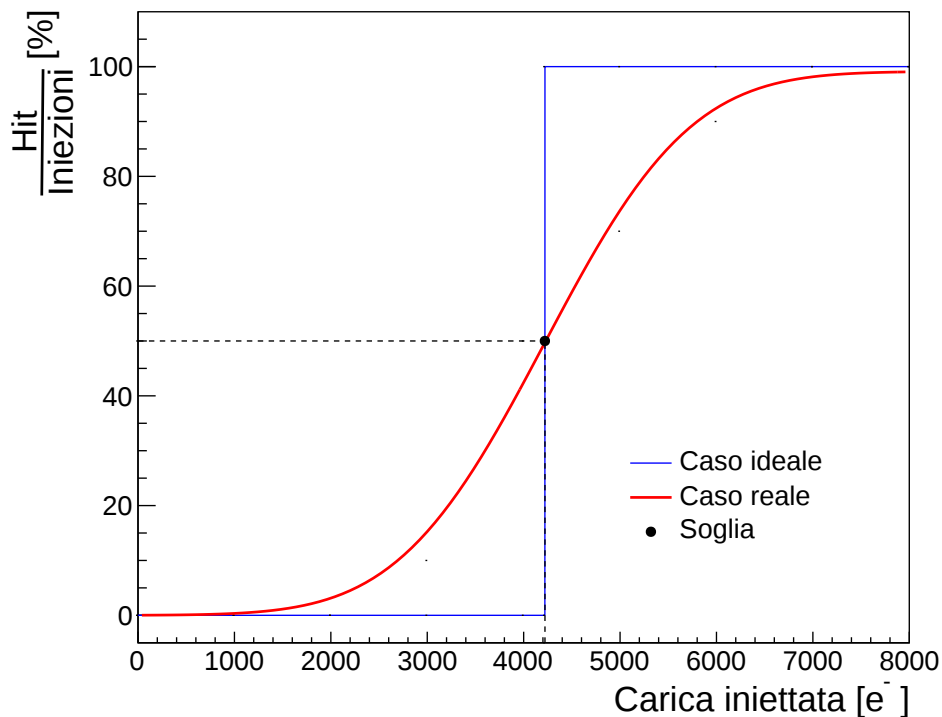


Figura 5.13: Esempio di curva S.

front-end lineare (a sinistra) e differenziale (a destra), eseguite con l'obiettivo di ottenere una soglia di 1500 e^- . A differenza di quanto schematizzato in figura 5.13 sull'asse y è riporta-

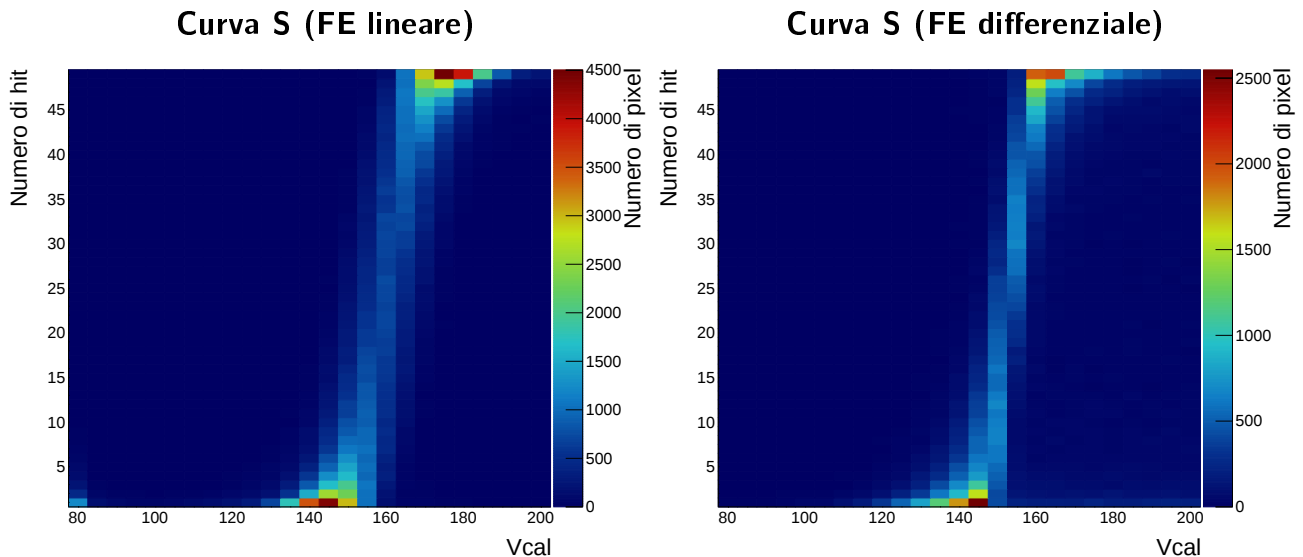


Figura 5.14: Curva S relativa al front-end lineare (sinistra) e differenziale (destra) del chip RD53A - SN: 0x1995 al termine della procedura di tuning fine della soglia locale.

to il numero degli hit registrati e non la loro percentuale rispetto al totale, mentre la carica iniettata Q è indicata sull'asse x in termini di $Vcal$, ovvero del parametro su cui agisce il ciclo Rd53aParameterLoop; la relazione tra le due è rappresentata da $Q \sim 10 \cdot Vcal$. In entrambi

i casi i pixel dei due front-end registrano il 50% delle iniezioni effettuate, pari cioè a 25, in corrispondenza di $V_{cal} \sim 150$, compatibilmente con quanto desiderato.

La procedura di tuning fine è più precisa rispetto a quella di tuning, ma più lenta; per questo motivo è eseguita effettuando un numero minore di iniezioni dopo aver già parzialmente ridotto la dispersione di soglia.

5.5.7 Scan della soglia

L'obiettivo di questa scansione è quello di misurare il valore della soglia impostato in ogni pixel.

I cicli eseguiti sono Rd53aMaskLoop, Rd53aCoreColLoop, Rd53aTriggerLoop e Rd53aParameterLoop, utilizzato per modificare il valore della carica iniettata. Al termine di ogni iterazione è realizzata una curva S per ciascun pixel; eseguendo un fit con una funzione degli errori è possibile ricavare il valore della soglia in ogni pixel, corrispondente alla media della funzione. Tale informazione è sfruttata per realizzare i grafici riportati in figura 5.15, ottenuti effettuando 50 iniezioni per ogni valore di carica. In alto a sinistra è rappresentata una mappa bidimensionale in cui a ciascun bin,

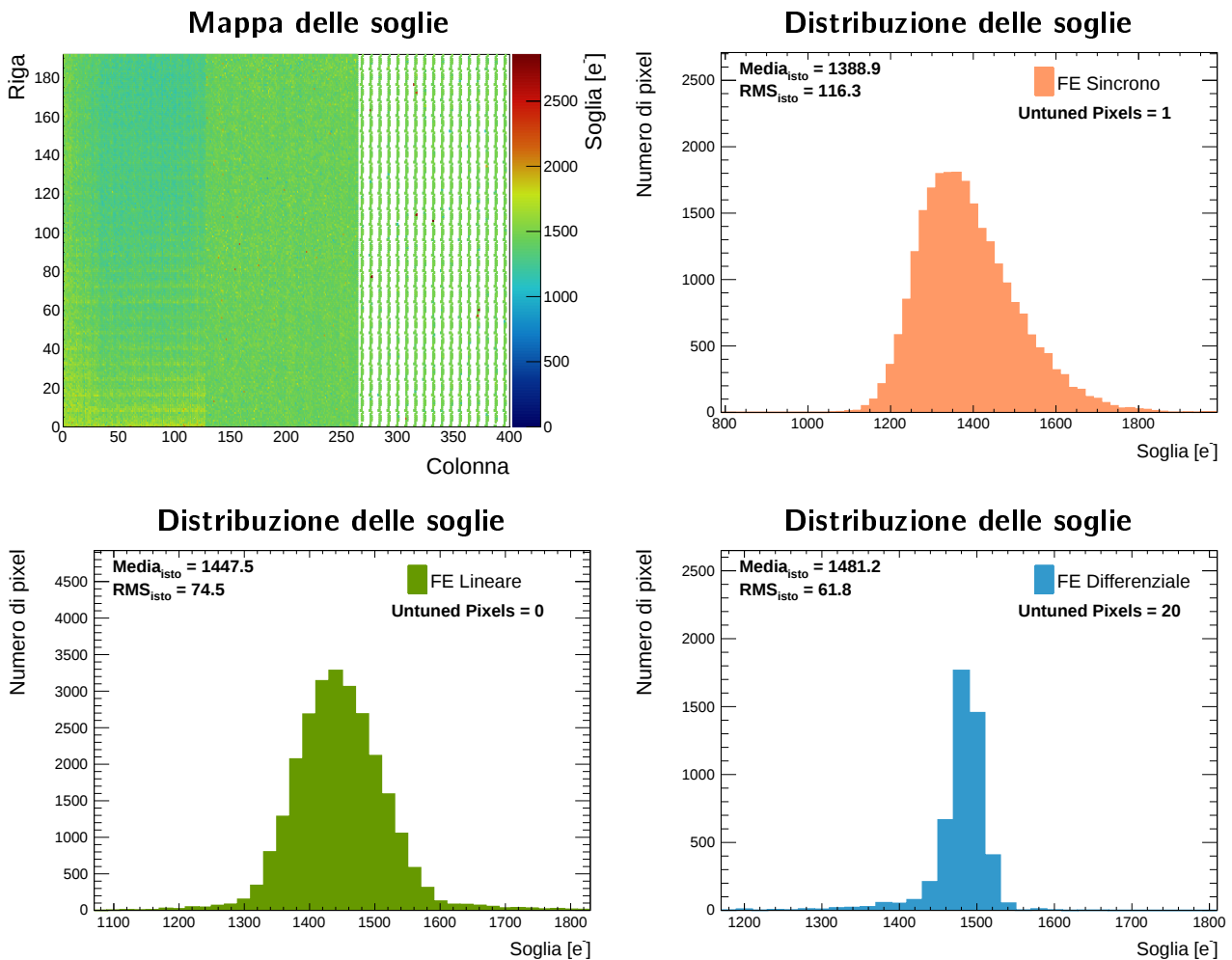


Figura 5.15: Mappa della soglia dei discriminatori nell'intera matrice dei pixel (in alto a sinistra), distribuzione relativa al front-end sincrono (in alto a destra), distribuzione relativa al front-end lineare (in basso a sinistra) e distribuzione relativa al front-end differenziale (in basso a destra) del chip RD53A - SN: 0x1995.

corrispondente ad un singolo pixel, è attribuito il valore della soglia; in alto a destra, in basso a

sinistra ed in basso a destra sono raffigurate le distribuzioni delle soglie relative rispettivamente al front-end sincrono, lineare e differenziale. In ogni grafico è possibile osservare i valori della media e dell'RMS dell'istogramma ed il numero di *untuned pixels*, ovvero dei pixel per i quali la procedura di tuning della soglia ha avuto esito negativo e pertanto essa ha valore nullo. Nella distribuzione relativa al front-end differenziale si osserva un minor numero di entrate; ciò è dovuto al fatto che anche in questo caso si sono considerati soltanto i valori delle soglie dei pixel abilitati dalla maschera differenziale. In questo esempio la soglia desiderata, indicata durante le procedure di tuning precedentemente descritte, è pari a 1500 e^- ; le medie ottenute sono compatibili, entro le dispersioni delle distribuzioni, con tale valore.

Sovrapponendo le curve S di ciascun pixel è possibile realizzare un grafico relativo all'intera matrice. In figura 5.16 è illustrato il risultato corrispondente ai grafici precedenti. Si osserva

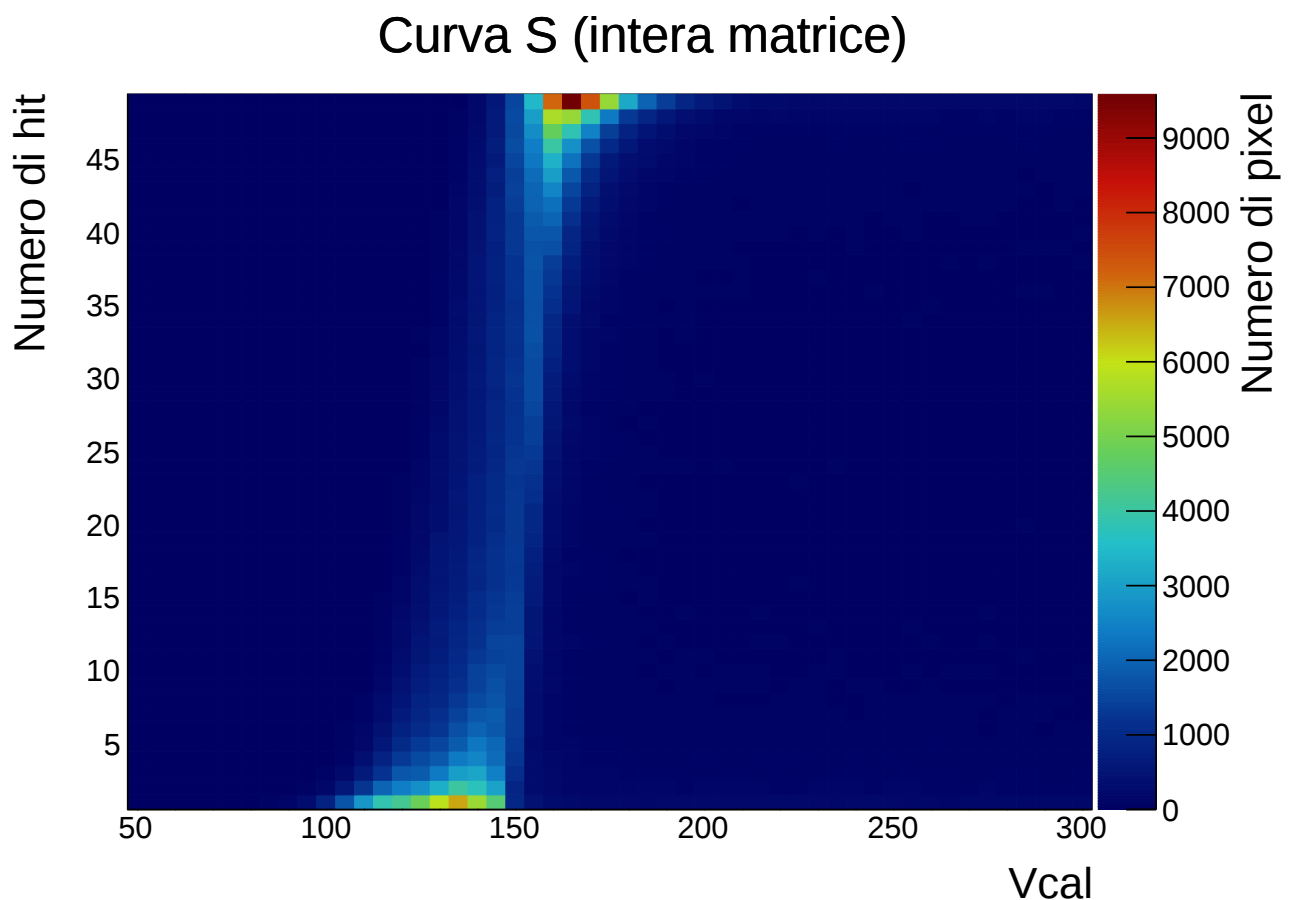


Figura 5.16: Esempio di curva S relativa all'intera matrice di pixel del chip RD53A - SN: 0x1995 ricavata al termine di uno scan di soglia in cui si sono effettuate 50 iniezioni in corrispondenza di ogni valore di carica.

come la curva S assuma il valore 25 (pari al 50% delle iniezioni effettuate) in corrispondenza di $V_{cal} \sim 150$, ovvero di una carica pari a $\sim 1500\text{ e}^-$, compatibilmente con quanto desiderato.

5.5.8 Scan del ToT

Questo scan è utilizzato per verificare l'uniformità di risposta dei pixel in termini di ToT in corrispondenza di un valore fissato di carica. Tale carica è iniettata 100 volte in ciascun pixel

della matrice del chip RD53A - SN: 0x1995, i quali effettuano il calcolo del ToT corrispondente ad ogni iniezione. Al termine della scansione a ciascun pixel è associato un valore di ToT, corrispondente alla media delle 100 misure effettuate.

In figura 5.17 sono illustrati i risultati di questo scan relativi all'iniezione di una carica pari a quella utilizzata nella procedura di tuning del CSA (sezione 5.5.5), ovvero 10000 e^- ; ci si aspetta dunque che a ciascun pixel sia associato un ToT medio pari a 8 bc. In alto a sinistra è

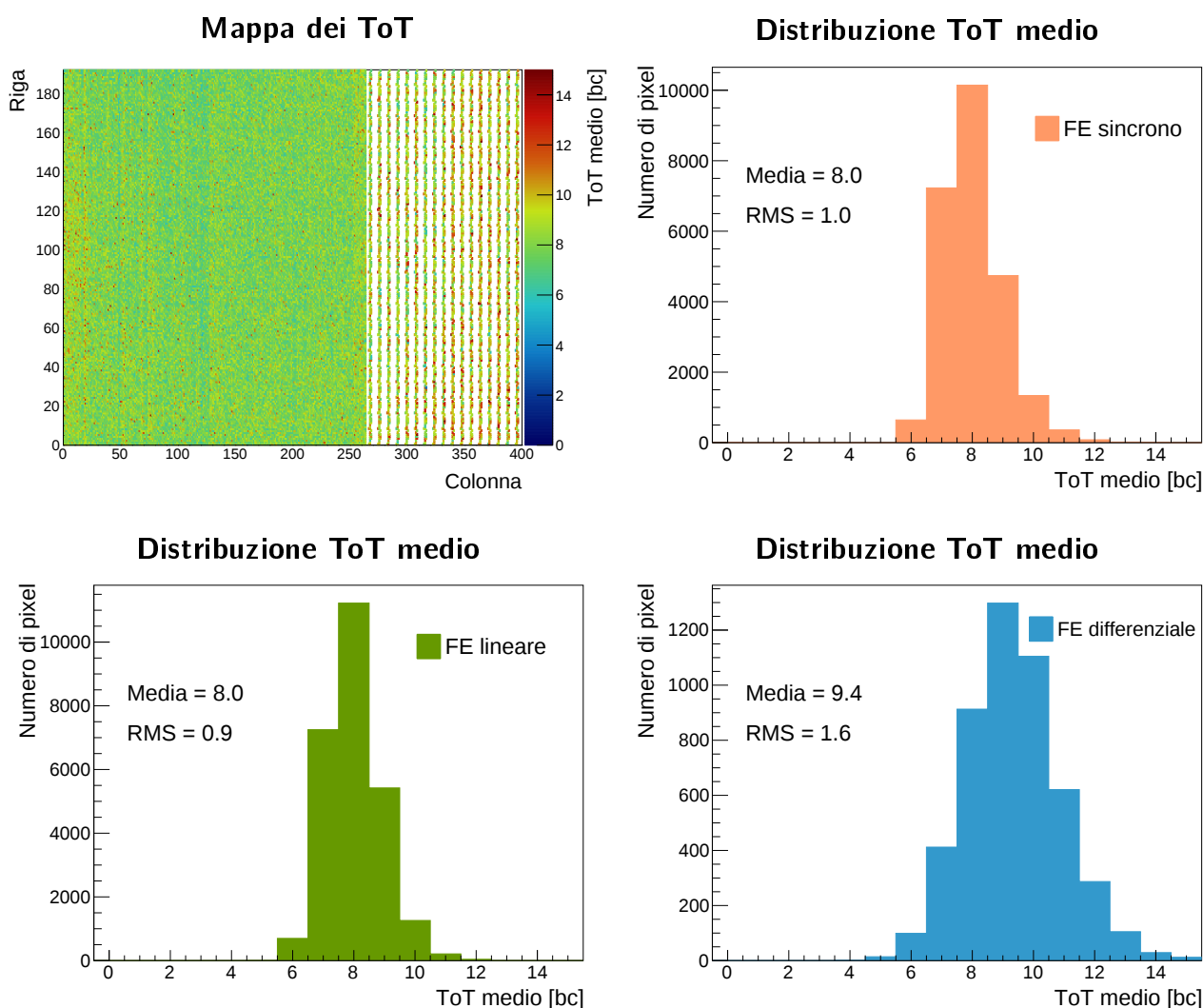


Figura 5.17: Risultato dello scan del ToT procedimento di tuning della corrente di scarica del CSA per il front-end sincrono (in alto a sinistra), lineare (in alto a destra) e differenziale (in basso) del chip RD53A - SN: 0x1995.

rappresentata la mappa bidimensionale contenente il valore medio dei ToT registrati in ciascun pixel, in alto a destra la distribuzione delle medie relative al front-end sincrono, in basso a sinistra quelle ottenute nel front-end lineare ed in basso a destra quelle del front-end differenziale. È possibile osservare come i valori medi siano vicini, entro le dispersioni delle distribuzioni, al valore atteso.

Capitolo 6

Misure di caratterizzazione del chip RD53A

La procedura suggerita per la corretta calibrazione del chip RD53A dagli sviluppatori del software di acquisizione dati YARR utilizzato nel presente lavoro di tesi è stata descritta nel capitolo 5, illustrando i risultati ottenuti per il chip RD53A - SN: 0x1995 in esame. In questo capitolo sono raccolte le misure effettuate con lo scopo di comprendere in maggiore dettaglio il comportamento dell'elettronica di front-end, in particolare dei DAC che generano la soglia di riferimento per il segnale prodotto da una particella ionizzante o, equivalentemente, indotto dall'utente, e dei circuiti che determinano la scarica dei condensatori integrati nei CSA e quindi la relazione di proporzionalità tra carica all'ingresso del front-end e ToT ad essa associato. Inoltre, è illustrato il risultato di una misura del rumore effettuata ed è stata valutata la qualità del sensore integrato sul chip, in termini di integrità dei bump-bond e di curva I-V. Tali analisi sono fondamentali per la corretta interpretazione dei risultati ottenuti esponendo il modulo ad una sorgente radioattiva, descritti nell'ultima sezione di questo capitolo.

6.1 Linearità dei DAC

La tensione di soglia inviata ad uno dei due ingressi dei discriminatori presenti in ogni circuito di front-end è generata a livello globale da convertitori digitale-analogico. I tre valori digitali inviati all'ingresso dei DAC, ciascuno relativo ad un diverso front-end, sono scritti nel file di configurazione del chip.

Una delle caratteristiche che definiscono le prestazioni di un DAC è la sua linearità, la quale indica la proporzionalità tra le variazioni del segnale digitale in ingresso e quelle del segnale analogico in uscita. In figura 6.1 è rappresentata la relazione esistente nel caso ideale tra il segnale in ingresso e quello in uscita da un DAC a 3 bit; il primo è espresso in termini di bit, il secondo in rapporto al massimo valore ottenibile. Il grafico in figura prende il nome di *funzione di trasferimento* del convertitore.

Per verificare la linearità dei DAC che forniscono la tensione di riferimento ai discriminatori del chip RD53A - SN: 0x1995 è possibile eseguire le procedure di tuning elencate nelle sezioni 5.5.3, 5.5.4, 5.5.5 e 5.5.6 al variare della soglia desiderata, indicata nel seguito come *soglia target*; tali scansioni modificano i valori dei segnali digitali all'ingresso dei DAC in modo da ottenere la soglia target. Effettuando uno scan della soglia, descritto nella sezione 5.5.7, si ottengono le distribuzioni illustrate in figura 5.15; per ciascuna di esse può essere eseguito un

Funzione di trasferimento DAC ideale

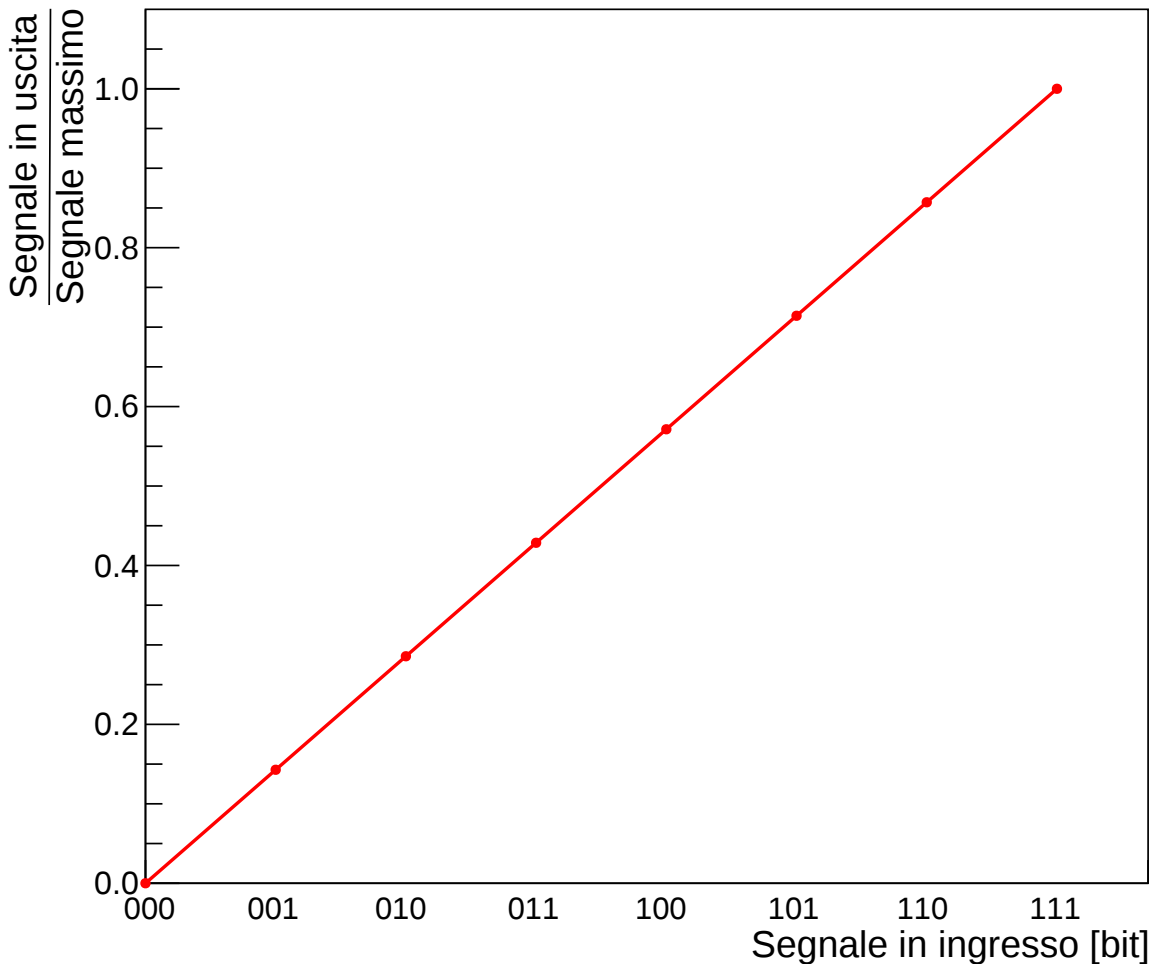


Figura 6.1: Funzione di trasferimento di un DAC a 3 bit nel caso ideale.

fit gaussiano, dal quale è possibile ricavare il valor medio della soglia relativo ad ogni front-end e la σ della distribuzione. Rappresentando tali medie in funzione dei valori digitali in ingresso ai DAC, salvati nel file di configurazione al termine delle procedure di tuning, si ricavano le funzioni di trasferimento delle quali si intende valutare la linearità.

In figura 6.2 sono illustrati i risultati ottenuti facendo variare la soglia target tra 900 e⁻ e 1900 e⁻ a passi di 100 e⁻. L'errore associato ad ogni valor medio del fit gaussiano è pari alla σ della distribuzione. In nero sono rappresentati i fit ai tre grafici, ottenuti utilizzando la relazione lineare $y = m \cdot x + q$; i valori risultanti da tale procedura, riportati in tabella 6.1, indicano la compatibilità dei tre fit con l'ipotesi lineare.

In corrispondenza di ciascun valore all'ingresso dei DAC è rappresentata, con un asterisco nero, la soglia target di ogni procedura di tuning; l'obiettivo è quello di verificare che questa sia compatibile con il valore medio risultante dal fit gaussiano della distribuzione della soglia, tenendo conto della sua dispersione. Si osserva che la procedura di tuning fine del software YARR converge ad un valore della soglia inferiore a quello della soglia target con una discrepanza che aumenta con l'aumentare di quest'ultima. Ciò indica che c'è margine per il miglioramento dell'algoritmo di tuning fine.

Funzione di trasferimento DAC globale

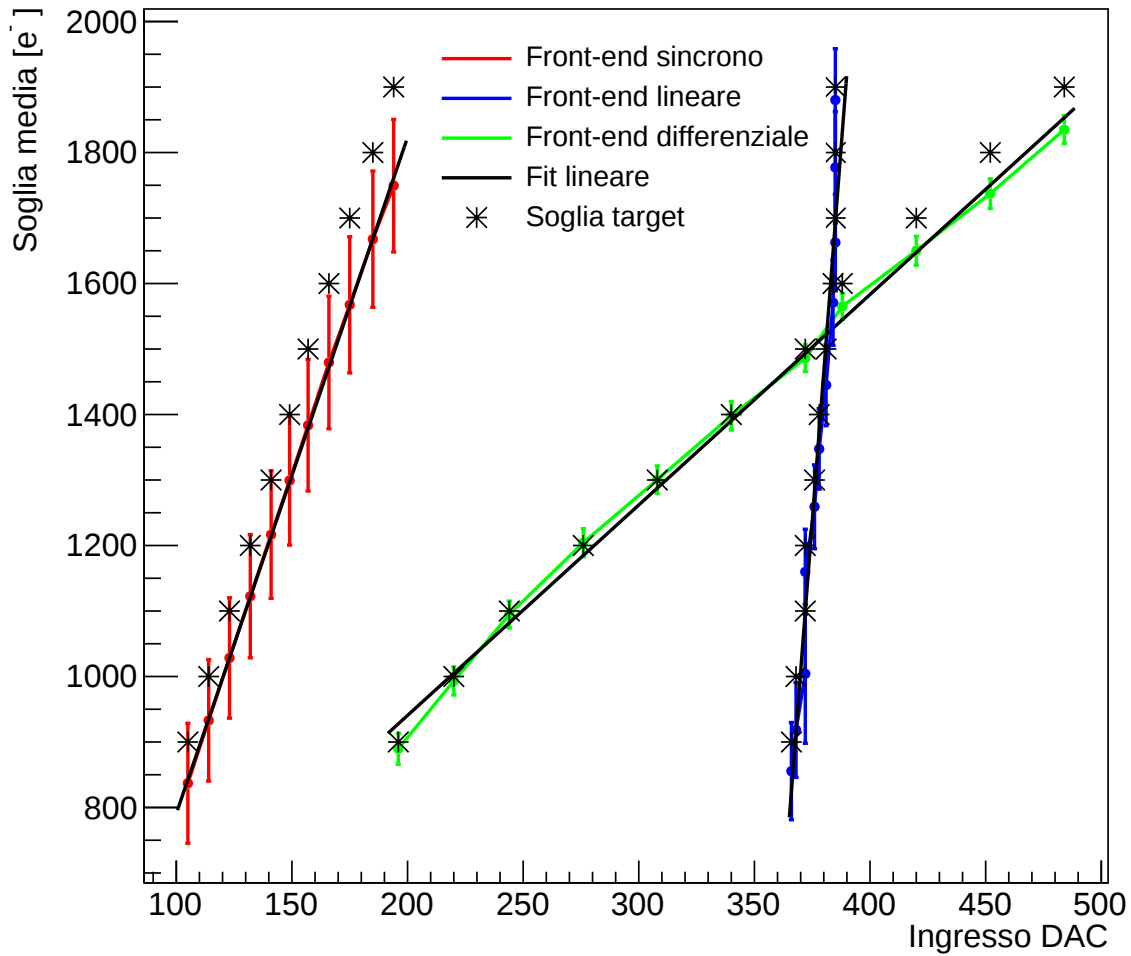


Figura 6.2: Fit lineare delle funzioni di trasferimento dei DAC che forniscono il valore di riferimento globale ai discriminatori contenuti nei tre diversi front-end del chip RD53A - SN: 0x1995.

	Front-end sincrono	Front-end lineare	Front-end differenziale
m	10.3 ± 1.1	45.7 ± 3.3	$(32.1 \pm 0.7) \cdot 10^{-1}$
q	$(-2.4 \pm 1.6) \cdot 10^2$	$(-15.9 \pm 1.2) \cdot 10^3$	$(29.8 \pm 2.5) \cdot 10$
χ^2/DOF	0.027/9	11.5/9	6.86/9
Prob.	1	0.24	0.65

Tabella 6.1: Valori risultanti dalla procedura di fit lineare delle funzioni di trasferimento dei DAC relativi ai tre diversi front-end, rappresentate in figura 6.2.

6.1.1 Isteresi della funzione di trasferimento

Le curve di trasferimento dei DAC potrebbero essere affette da problemi di isteresi; ovvero, la risposta del chip ad un determinato tuning potrebbe essere influenzata dal valore della soglia target del tuning precedente.

Per verificare tale eventualità la procedura descritta nella sezione 6.1 è stata ripetuta in ordine inverso, diminuendo la soglia target da 1900 e^- a 900 e^- . Il confronto tra i risultati

ottenuti è rappresentato in figura 6.3, nella quale le linee tratteggiate sono relative al caso in cui la soglia è stata progressivamente diminuita. È possibile osservare un discostamento tra la

Isteresi funzione di trasferimento DAC globale

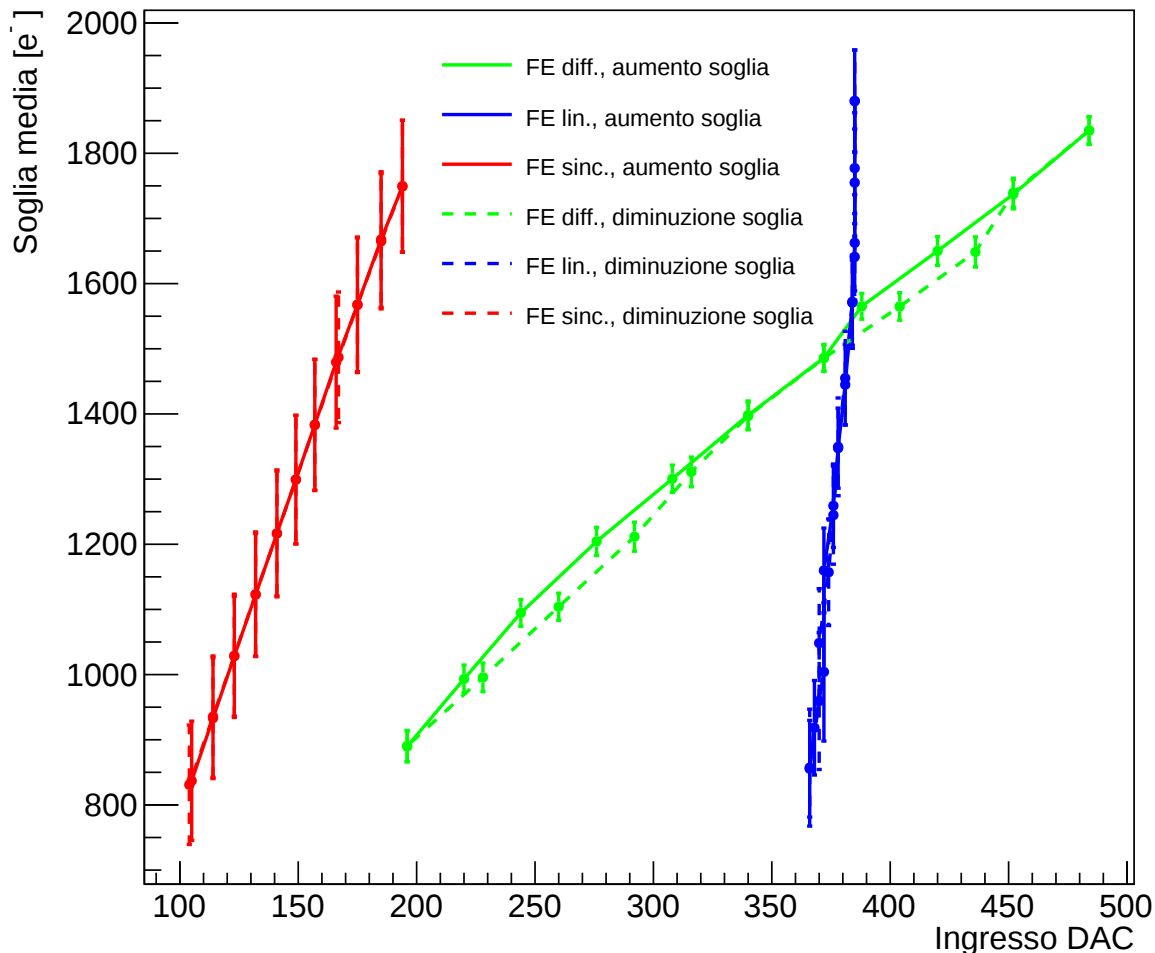


Figura 6.3: *Isteresi nelle funzioni di trasferimento dei DAC globali dei tre front-end.*

curva continua e quella tratteggiata solo nel caso del front-end differenziale. I valori medi delle distribuzioni relative allo stesso valore di soglia target sono comunque compatibili tra loro, per cui si può concludere che il fenomeno non è rilevante.

6.2 Impatto del tuning locale della soglia

Il tuning locale della soglia, descritto nella sezione 5.5.4, è eseguito per ridurre la dispersione di soglia presente nel front-end differenziale ed in quello lineare al termine della procedura di tuning globale.

Per verificare l'esito di tale scansione è possibile effettuare un confronto tra le larghezze delle distribuzioni di soglia (misurate come σ dei rispettivi fit gaussiani) relative alla stessa soglia target al termine della procedura di tuning globale e di quella locale. In figura 6.4 è rappresentato il risultato ottenuto al variare della soglia target tra 900 e^- e 1900 e^- a passi di 100 e^- ;

Dispersione soglia

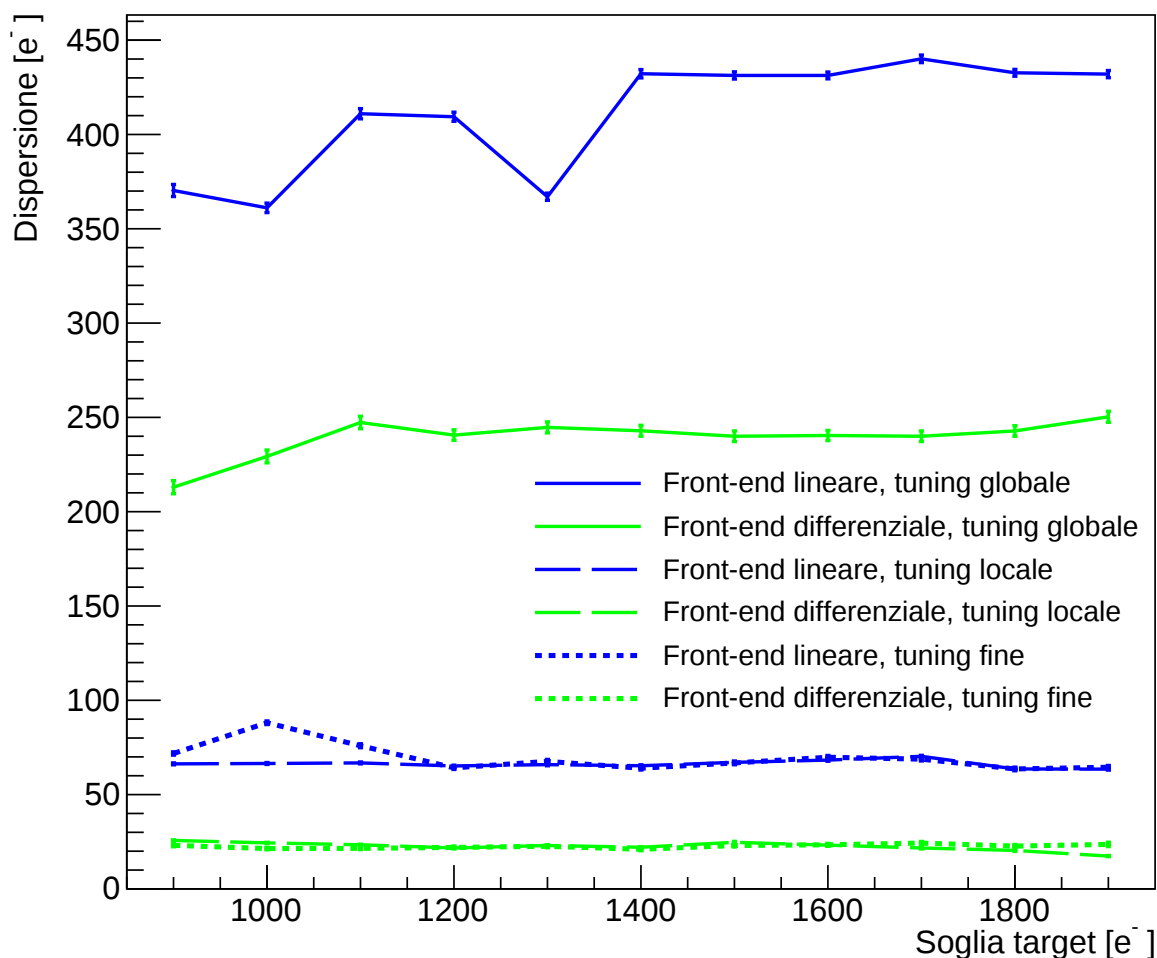


Figura 6.4: Confronto tra la dispersione di soglia ottenuta al termine della procedura di tuning globale (linea continua), locale (linea tratteggiata) e di tuning fine (linea a punti) per i front-end lineare (in blu) e differenziale (in verde) del chip RD53A - SN: 0x1995.

l'incertezza associata ad ogni misura è pari all'errore sulla σ del fit gaussiano. Le curve relative al tuning globale sono indicate con una linea continua, quelle relative al tuning locale con una linea tratteggiata. Il front-end sincrono, per il quale non esiste una procedura di tuning locale, è escluso da tale analisi. È evidente come la dispersione di soglia diminuisca notevolmente in seguito alla procedura di tuning locale; si può quindi concludere che tale scansione è fondamentale per configurare il chip RD53A - SN: 0x1995 in modo ottimale.

Come spiegato nella sezione 5.5.6, la procedura di tuning fine permette di ridurre la dispersione di soglia in modo più efficace, sebbene più lento, rispetto al semplice tuning locale. Per verificare tale ipotesi in figura 6.4 è illustrata anche la σ della distribuzione di soglia ricavata al termine della procedura di tuning fine per ogni valore di soglia target. La curva ottenuta, indicata con una linea a punti, si discosta in maniera impercettibile da quella relativa al tuning locale (ad eccezione del punto in corrispondenza di una soglia target di 1000 e⁻). La procedura di tuning fine, quindi, non appare indispensabile per la corretta configurazione del chip in esame.

6.3 Influenza della corrente di scarica sul ToT

Come illustrato in figura 4.7 la corrente di scarica del condensatore integrato nel CSA influenza la velocità di ritorno alla baseline del segnale in uscita dall'amplificatore, e dunque il valore del ToT corrispondente ad una certa carica all'ingresso del front-end. Per verificare tale effetto si esegue un tuning del CSA in modo che una carica di 10000 e^- corrisponda ad un ToT di 8 bc e successivamente si effettuano 3 diversi scan del ToT aumentando progressivamente di 20 unità i parametri che determinano la scarica del condensatore e quindi la velocità di ritorno alla baseline del segnale. I risultati ottenuti sono riportati in figura 6.5; si osserva come i valori medi di tale distribuzioni diminuiscano procedendo dall'alto verso il basso, ovvero all'aumentare dei parametri in esame. Le distribuzioni in figura sono dunque in accordo con le previsioni teoriche, illustrate in figura 4.7.

6.4 Linearità del ToT

Come spiegato nella sezione 5.5.5, la procedura di tuning del CSA consente di modificare la proporzionalità esistente tra la carica all'ingresso dell'amplificatore ed il valore di ToT in uscita dal discriminatore ad essa corrispondente. Per verificare la linearità tra la misura di ToT e la carica all'ingresso del CSA, dopo aver fissato ed equalizzato la calibrazione dei CSA del chip, è possibile eseguire una successione di scan del ToT al variare della quantità di carica iniettata Q e calcolare il valor medio delle distribuzioni delle medie dei ToT registrati all'interno di ciascun front-end, analoghe a quelle rappresentate in figura 5.12.

Il tuning del CSA è effettuato in modo tale che una carica di 10000 e^- all'ingresso del front-end corrisponda ad un ToT medio di 8 bc. Successivamente si eseguono 30 ripetizioni della misura del ToT per ogni pixel variando Q tra 1200 e^- e 30000 e^- , che rappresentano gli estremi dell'intervallo consentito dal circuito di calibrazione. I risultati ottenuti sono rappresentati nelle tre mappe illustrate in figura 6.6, ciascuna relativa ad un diverso front-end, nelle quali i valori di Q sono rappresentati sull'asse x e le distribuzioni delle medie dei ToT sull'asse y ; sull'asse z è indicata la legenda dei colori associati al numero di pixel relativo a ciascun bin dell'asse y . Il front-end sincrono e quello lineare manifestano il comportamento atteso: il valor medio delle 30 misure ha un andamento lineare all'aumentare di Q ed è pari a 8 bc in corrispondenza di $Q = 10000\text{ e}^-$. La scelta di far corrispondere tale valore di carica con la metà del massimo ToT ottenibile (si ricorda che il calcolo del ToT è effettuato da un contatore a 4 bit) consente di verificare anche la saturazione delle due curve, visibile per valori di Q superiori a $\sim 24000\text{ e}^-$.

Nel front-end differenziale il minor numero di pixel analizzati è dovuto all'applicazione della maschera differenziale. Il valor medio corrispondente a $Q = 10000\text{ e}^-$ è pari a 9 bc; esso è comunque compatibile con il valore atteso considerando la larghezza della distribuzione. Tuttavia, la relazione tra ToT medio e Q in questo front-end si discosta dall'andamento lineare e il ToT medio non raggiunge mai il valore massimo di fondo scala, pari a 15 bc. Questo risultato conferma quello ottenuto da altri gruppi appartenenti alla collaborazione ATLAS che hanno svolto analoghe misure su altri chip RD53A.

Distribuzioni ToT medio

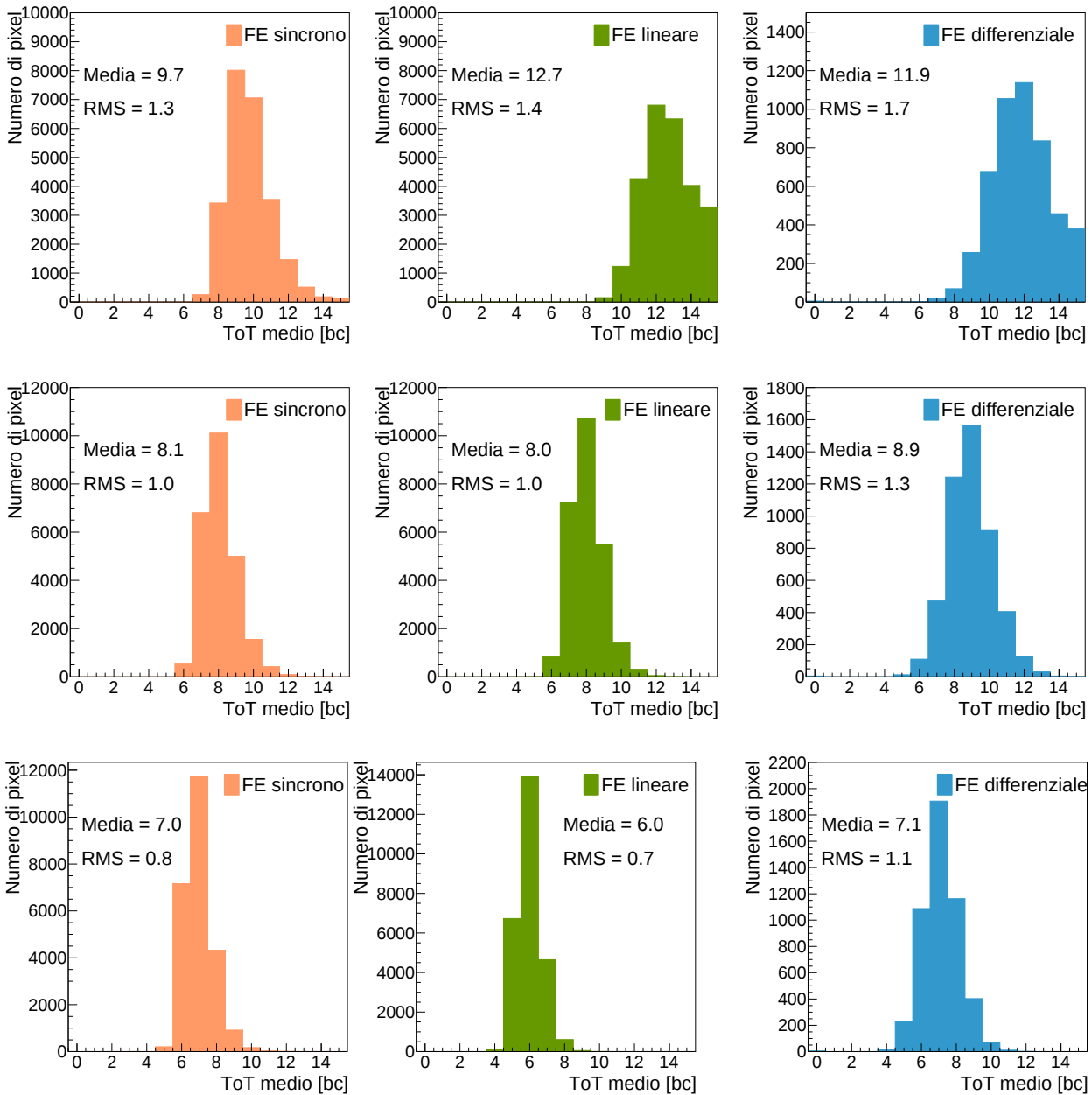


Figura 6.5: Distribuzioni dei valori medi di ToT ottenute aumentando progressivamente i parametri che definiscono la scarica dei condensatori integrati nei CSA del chip RD53A - SN: 0x1995, e quindi la velocità di ritorno alla baseline del segnale.

6.5 Scan del rumore

L'inevitabile presenza di rumore elettronico può produrre segnali di ampiezza maggiore rispetto alla tensione di soglia; in tal caso è registrato un hit non corrispondente al rilascio di energia da parte di una particella o ad una iniezione di carica. Per individuare i pixel nei quali il rumore è particolarmente elevato è possibile effettuare uno scan del rumore.

In questa procedura di test l'unico ciclo eseguito tra quelli elencati nella sezione 5.4.1 è Rd53aTriggerLoop, utilizzato per la sola generazione dei segnali di trigger: l'iniezione di carica

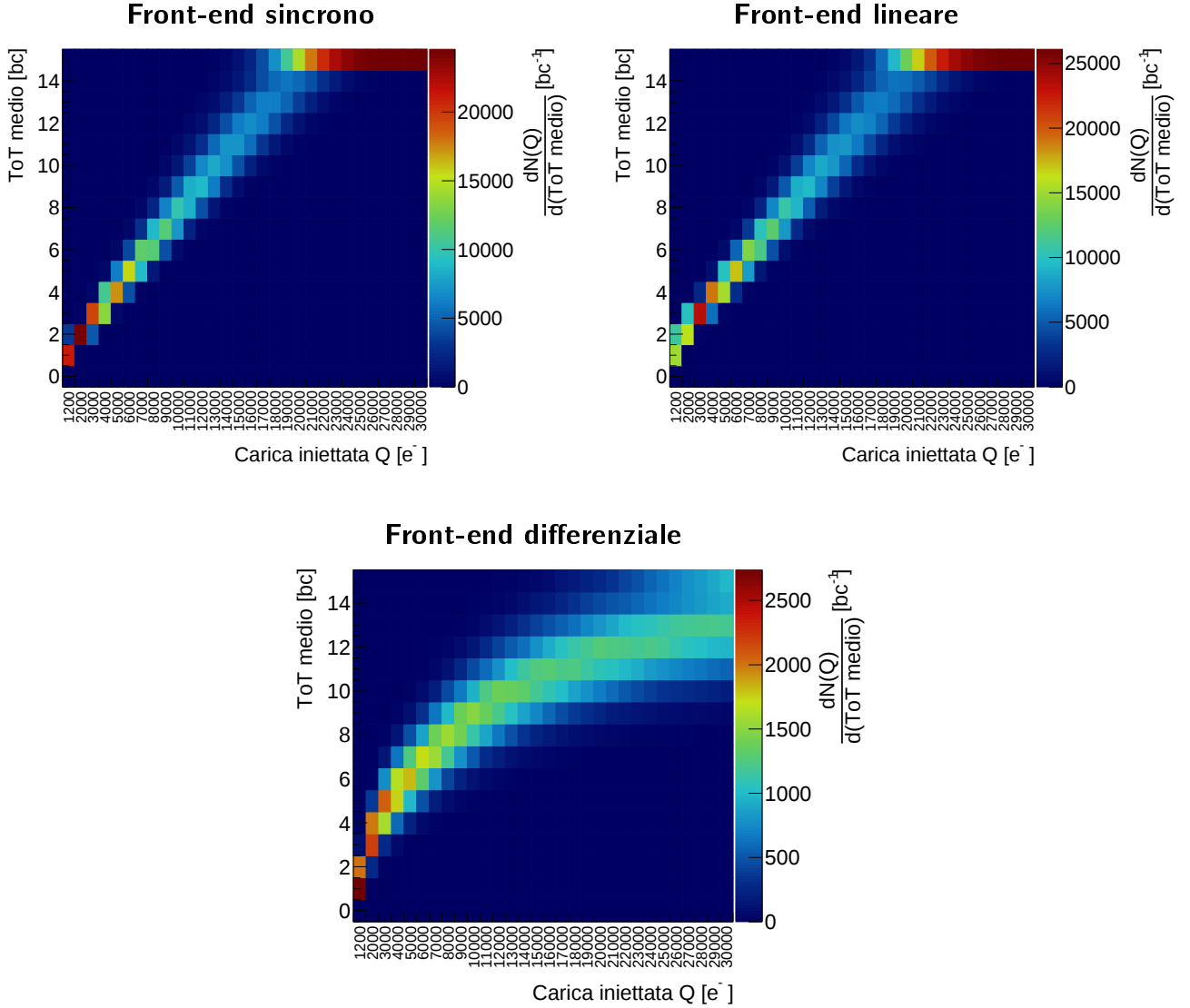


Figura 6.6: Mappe rappresentanti le distribuzioni dei valori medi di ToT (in y) corrispondenti a diverse cariche iniettate (in x), relative al front-end sincrono (in alto a sinistra), lineare (in alto a destra) e differenziale (in basso) del chip RD53A - SN: 0x1995.

è disabilitata e mentre la frequenza di generazione del trigger e la durata della presa dati sono definite dall'utente. In tal modo l'intera matrice di pixel è letta ad intervalli di tempo regolari per verificare l'eventuale presenza di segnali sopra soglia i quali, non essendo stata effettuata alcuna iniezione, sono necessariamente generati dal rumore.

I risultati ottenuti al termine di uno scan della durata di 5 minuti eseguito dopo aver configurato il chip in modo da regolare la soglia a $1000 e^-$ sono illustrati in figura 6.7. Nel grafico a sinistra è rappresentato il numero di hit registrati da ciascun pixel per unità di bc; nella mappa a destra è assegnato valore 0 ai pixel per i quali la quantità rappresentata a sinistra è superiore a 10^{-6} hit/bc. Questi possono dunque essere facilmente individuati e disabilitati durante la fase di operazione del chip.

Il valore 10^{-6} hit/bc rappresenta il limite superiore imposto ai pixel dei chip di lettura che saranno utilizzati nell'ITk [48]. Poiché RD53A contiene 76800 pixel, ciò implica che nell'intera matrice si può tollerare una frequenza massima di segnali generati dal rumore pari a 0.08 hit/bc;

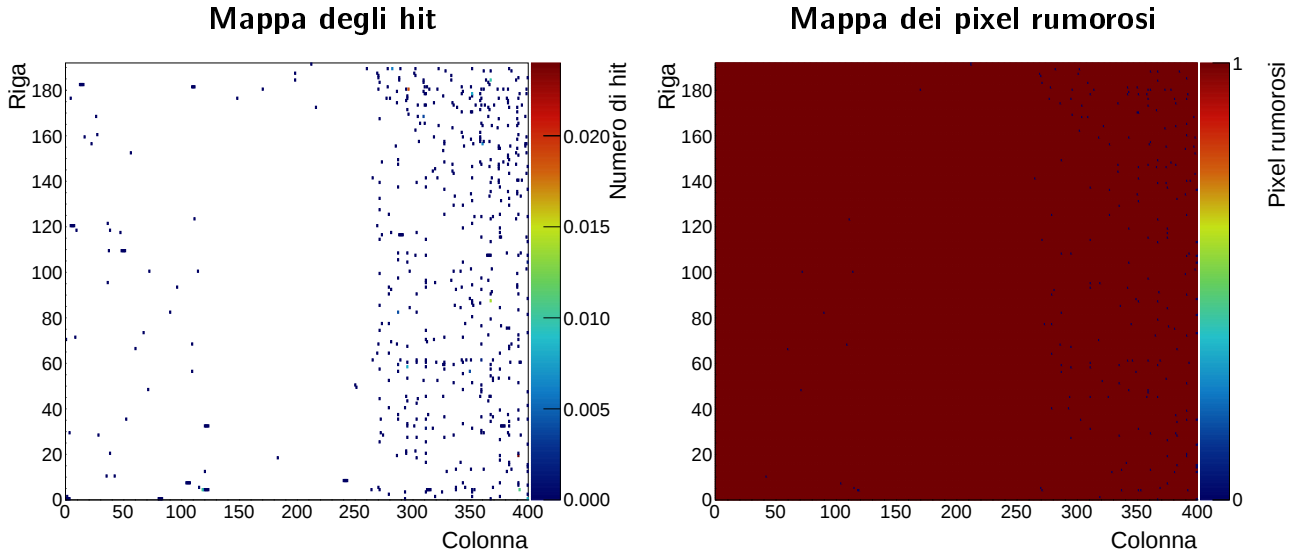


Figura 6.7: Mappe raffiguranti il numero di hit registrati da ogni pixel del chip RD53A - SN: 0x1995 in unità di bc (a sinistra) ed i pixel per i quali tale quantità supera 10^{-6} hit/bc, associati al valore 0 (a destra).

tale valore è notevolmente inferiore rispetto a quello di hit reali previsto nei chip utilizzati nello strato più interno (dell'ordine delle centinaia di hit/bc) ed in quelli successivi di ITk (dell'ordine delle decine di hit/bc).

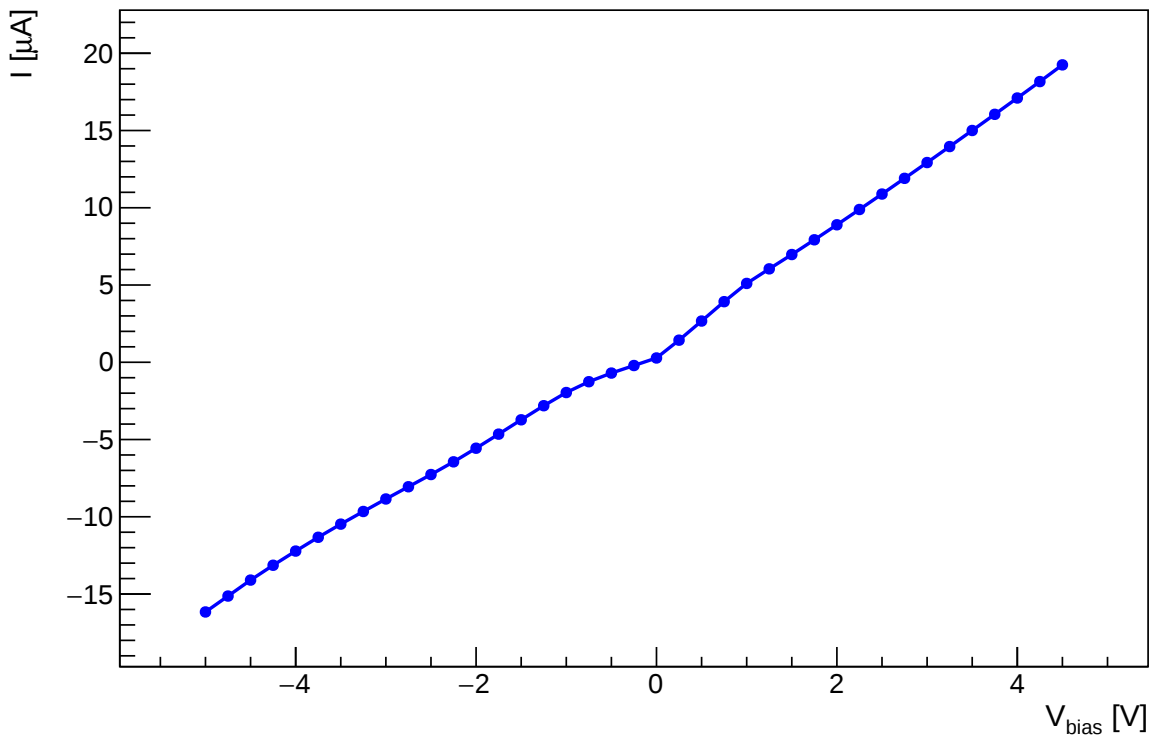
6.6 Curva I-V

Sul chip RD53A - SN: 0x1995 in esame è integrato, tramite bump-bonding, un sensore 3D. Le sue prestazioni possono essere valutate analizzando la sua curva I-V, introdotta nella sezione 3.2.1. Per realizzare tale curva si varia la tensione di alimentazione del sensore V_{bias} tra 4.5 V e -5 V e si misura il valore della corrente di deriva in corrispondenza di ogni decremento della tensione, con passi pari a 0.25 V. L'alimentazione è fornita dal generatore Keithley 6487, utilizzato anche per fornire la lettura della corrente; per ridurre al minimo il rumore dovuto alla radiazione luminosa ambientale si copre il modulo con un telo nero. Il risultato ottenuto è rappresentato in figura 6.8.

Il comportamento di un sensore ideale è quello illustrato in figura 3.3. Nella regione di polarizzazione inversa in cui $V_{bias} < 0$, ovvero quella nella quale è adoperato il rivelatore, non si osserva l'andamento $I \propto \sqrt{|V_{bias}|}$ atteso; inoltre, la risposta del sensore nell'intero intervallo di tensione considerato può essere assimilata a quella di un dispositivo ohmico⁽¹⁾ ed è pertanto significativamente differente da quella rappresentata in figura 3.3. Si può quindi concludere che il sensore analizzato non funziona correttamente.

⁽¹⁾Dispositivo per il quale vale la legge di Ohm $V = RI$, in cui R indica la sua resistenza, V la differenza di potenziale ad esso applicata e I la corrente generata al suo interno.

Curva I-V

**Figura 6.8:** Curva I-V del sensore integrato sul chip RD53A - SN: 0x1995.

6.7 Integrità dei bump-bond

Come introdotto nella sezione 3.5.1 la realizzazione dei bump-bond presenta notevoli difficoltà; per questo motivo è opportuno verificare che il sensore sia correttamente connesso al chip RD53A - SN: 0x1995 in esame.

Un modo per valutare lo stato dei bump-bond consiste nel verificare la presenza di *cross-talk* nel chip. Con questo termine si indica il fenomeno per il quale la carica rilasciata in un pixel induce un hit anche in quelli adiacenti a causa di capacità parassite tra di essi. Dal momento che tali capacità sono inevitabilmente presenti nel modulo, l'assenza di cross-talk in un pixel è interpretata come un danneggiamento del bump-bond ad esso corrispondente e dunque una non connessione al sensore. Tale conclusione si basa sul fatto che il cross-talk indotto dall'elettronica del chip sia trascurabile e che quindi tale fenomeno sia attribuibile unicamente al sensore.

Per verificare l'esistenza di cross-talk è possibile contare il numero di hit rivelati in un pixel in corrispondenza dell'iniezione di una certa quantità di carica in quelli adiacenti, come rappresentato in figura 6.9: in celeste sono indicati i pixel nei quali avviene l'iniezione di carica, in arancione quello in cui si conta il numero di hit registrati. In assenza di cross-talk tale numero è pari a 0.

La tensione fornita al generatore per effettuare questa verifica è pari a -2 V, mentre la soglia di discriminazione è fissata a $1100 e^-$. A sinistra in figura 6.10 è rappresentata la mappa degli hit ottenuta al termine di uno scan in cui si sono effettuate 100 iniezioni di carica nei pixel selezionati dai cicli Rd53aMaskLoop e Rd53aCoreColLoop, definiti in modo da iniettare carica in un pattern di pixel contigui a quello sottoposto al test, secondo lo schema in figura 6.9 (al termine

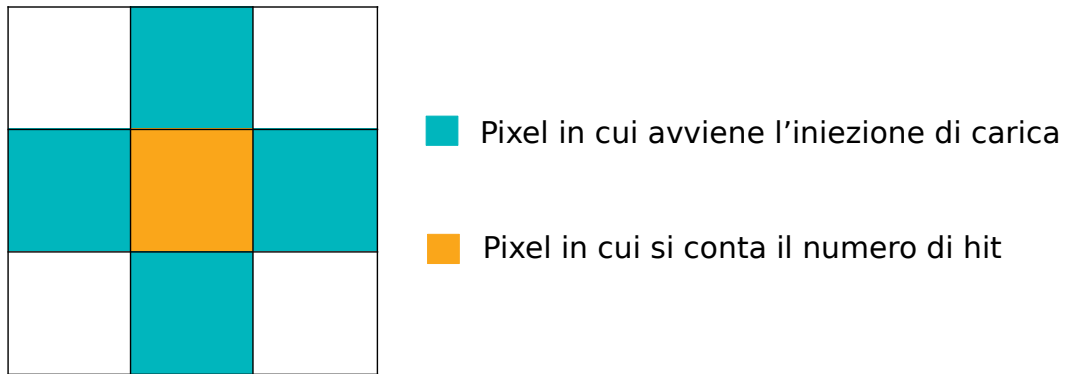


Figura 6.9: Schematizzazione del metodo utilizzato per verificare la presenza di cross-talk nel pixel centrale.

della scansione ogni pixel del chip risulta sottoposto al test di cross-talk). È possibile osservare

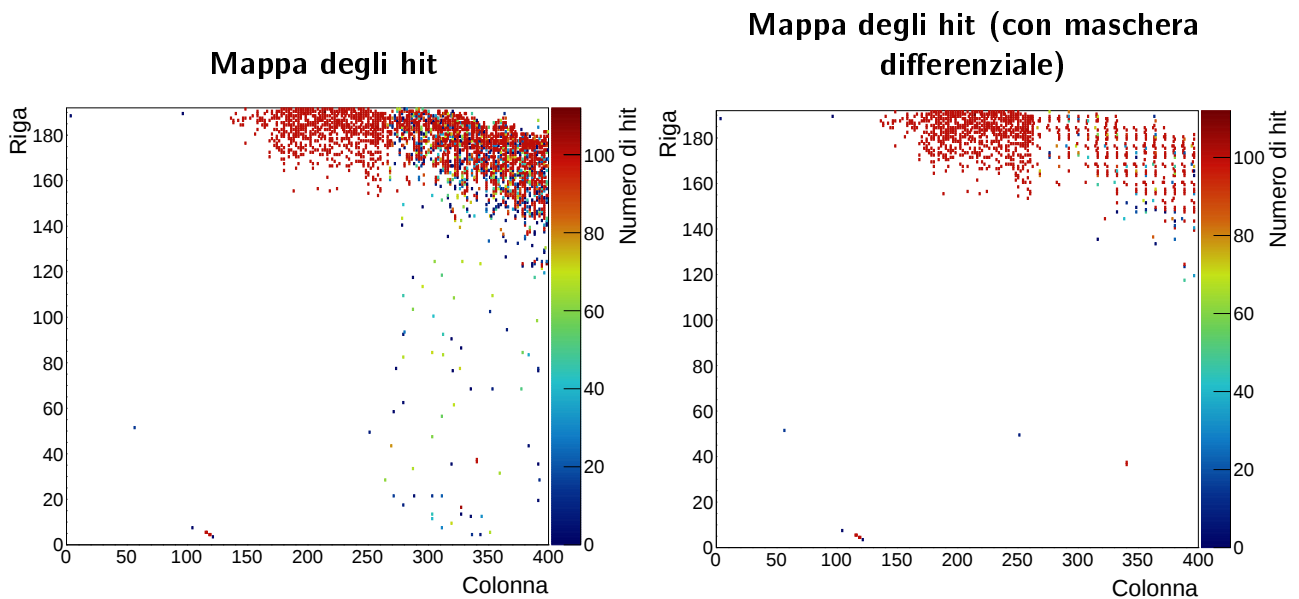


Figura 6.10: Sinistra: Mappa degli hit registrati in ciascun pixel in corrispondenza dell'iniezione di carica nei pixel adiacenti con l'obiettivo di verificare la presenza di cross-talk nel modulo costituito dal chip RD53A - SN: 0x1995 e dal sensore ad esso collegato. Destra: riproduzione della mappa a sinistra con l'applicazione della maschera differenziale.

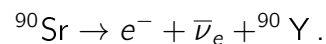
come nella maggior parte dei pixel non si sia evidenziato alcun cross-talk. Il front-end sincrono (si ricorda che esso è implementato nelle colonne di pixel 0-127) è completamente disconnesso dal sensore, mentre nel front-end lineare (implementato nelle colonne 128-255) l'unica regione effettivamente connessa al sensore è quella collocata in alto (con riferimento alla figura). Il front-end nel quale si osserva una maggiore integrità nei bump-bond è quello differenziale (implementato nelle colonne di pixel 256-400); per riferimento, a destra in figura 6.10 è illustrata la stessa mappa ottenuta dopo l'applicazione della maschera differenziale riportata in figura 5.10.

Il risultato illustrato indica che il modulo in esame è gravemente danneggiato e che una frazione consistente dei bump-bond che collegano correttamente il sensore al front-end differenziale è integrata su pixel che forniscono una risposta scorretta; ciò deve essere tenuto in considerazione nelle analisi condotte utilizzando il sensore. Dal momento che questo è collegato a massa tramite i bump-bond, il fatto che la maggior parte di essi risulti danneggiata potrebbe spiegare il malfunzionamento del sensore indicato dalla curva I-V illustrata nella sezione 6.6.

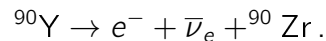
6.8 Esposizione ad una sorgente radioattiva

Il funzionamento del chip RD53A - SN: 0x1995 è stato verificato sino ad ora iniettando all'ingresso dei circuiti di front-end in esso integrati una quantità di carica stabilita dall'utente, o rivelando gli hit prodotti dal rumore elettronico. Per indagare il suo comportamento in risposta al segnale generato dall'interazione con una particella, il modulo costituito da tale chip e dal sensore 3D ad esso collegato è stato esposto ad una sorgente radioattiva.

La sorgente scelta è ^{90}Sr ; esso è un isotopo radioattivo dello Stronzio soggetto ad un decadimento β con un'emivita di 28.79 anni:



L'isotopo dell'Ittrio ^{90}Y , anch'esso instabile, decade β con un'emivita di 64 ore secondo la reazione:



Poiché ^{90}Zr è un isotopo stabile dello Zirconio, tale decadimento rappresenta l'ultimo tra i processi innescati da una sorgente di ^{90}Sr , la quale è quindi una fonte di particelle β provenienti da due decadimenti consecutivi.

Il modulo in esame è stato calibrato regolando la soglia a 1100 e^- e la relazione di proporzionalità tra carica in ingresso e ToT in modo che 10000 e^- corrispondano ad un ToT di 8 bc, mentre il sensore è stato alimentato con una tensione pari a -2 V. Inoltre, si è eseguito uno scan del rumore con l'obiettivo di individuare e disabilitare i pixel che, in assenza di iniezione di carica, registrano un numero di hit superiore alla soglia di 10^{-6} hit/bc. Al termine di questa fase di calibrazione il chip è stato esposto alla sorgente di ^{90}Sr per un tempo pari a 5 minuti. La presenza di segnali è ricercata eseguendo uno scan del rumore: la lettura di tutti i pixel della matrice, la quale avviene in una finestra temporale corrispondente a 64 bc, è abilitata dalla ricezione di un segnale di trigger, inviato con una frequenza di 5 kHz. I dati raccolti sono stati utilizzati per realizzare la mappa degli hit registrati da ciascun pixel, illustrata in figura 6.11; nel front-end differenziale si sono considerati solo gli hit registrati dai pixel abilitati dalla maschera in figura 5.10. Si osserva che non tutti i pixel che hanno rivelato un segnale erano stati indicati come correttamente collegati al chip RD53A - SN: 0x1995 in figura 6.10; ciò può essere attribuito al fatto che il segnale generato dal cross-talk in tali pixel durante la procedura illustrata nella sezione 6.7 fosse di ampiezza inferiore rispetto alla soglia di discriminazione. Il numero medio di hit registrato dai pixel che hanno individuato un segnale prodotto dagli elettroni emessi dalla sorgente è pari a 8; per stabilire la validità di tale risultato è possibile stimare il valore

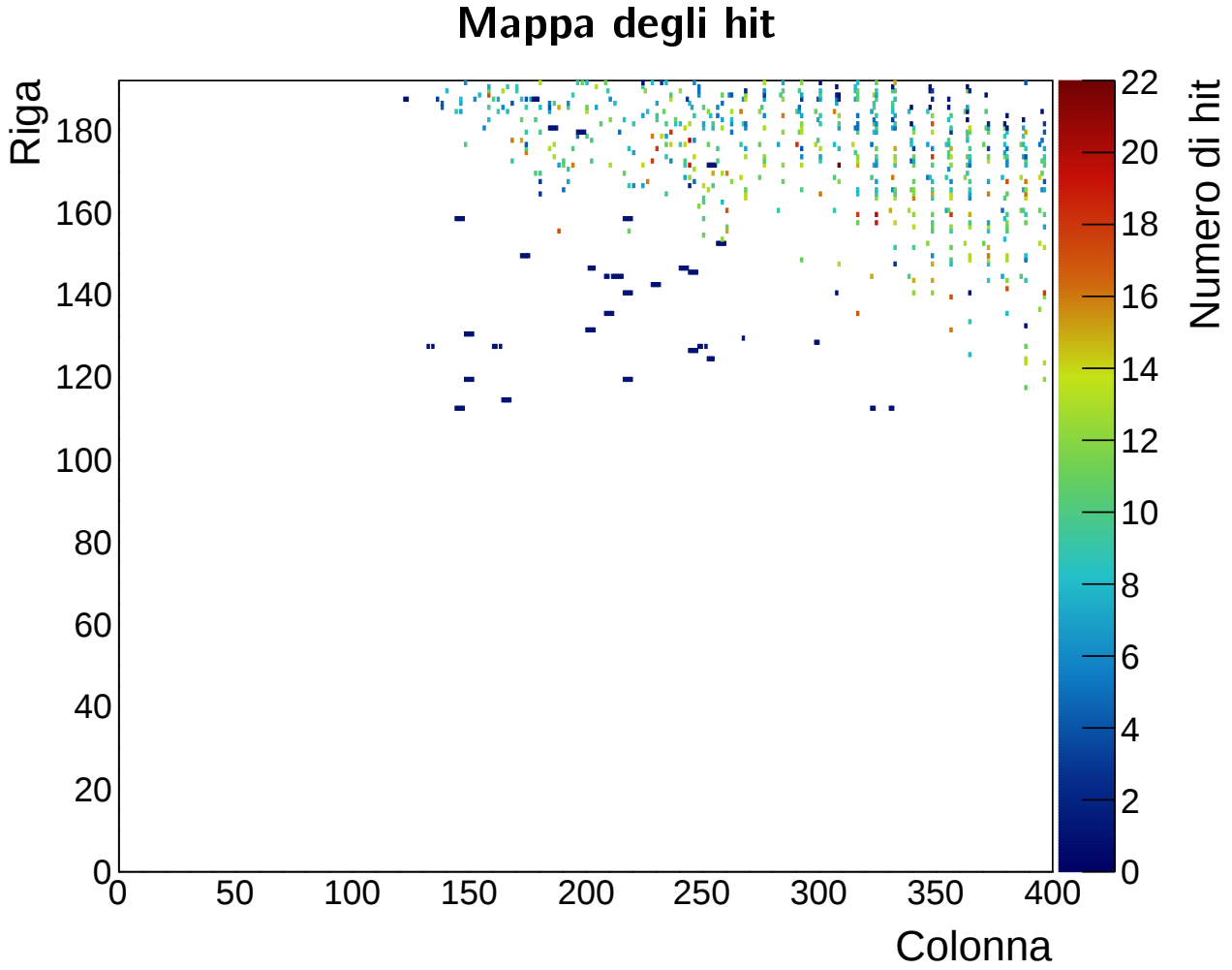


Figura 6.11: Mappa degli hit registrati da ciascun pixel durante l'esposizione del modulo ad una sorgente di ^{90}Sr .

atteso. La sorgente di ^{90}Sr adoperata ha un'attività di 3.7 MBq, che equivale alla produzione di 3.7×10^6 elettroni ogni secondo. Questi sono emessi su tutto l'angolo solido ed interagiscono con il sensore posto ad una distanza $r = 3$ cm dalla sorgente. Pertanto, il numero di elettroni attesi in ogni secondo su ciascun pixel, di superficie $A = 25 \mu\text{m} \times 100 \mu\text{m}$, è pari al prodotto tra la frequenza di emissione per unità di angolo solido, f , e l'angolo solido sotteso da ciascun pixel, Ω :

$$f \cdot \Omega = f \cdot \frac{A}{r^2} = \frac{3.7 \times 10^6 \text{ e}^- \cdot \text{s}^{-1}}{4\pi} \cdot \frac{25 \mu\text{m} \times 100 \mu\text{m}}{(3 \text{ cm})^2} = 0.8 \text{ e}^- \cdot \text{s}^{-1}. \quad (6.1)$$

Per poter stimare il numero di hit attesi in ciascun pixel al termine dello scan eseguito è necessario calcolare il tempo effettivo di acquisizione dati. I segnali di trigger sono inviati al chip con una frequenza di 5 kHz, che equivale ad un periodo di 200 μs . Il modulo è stato esposto alla sorgente per 5 minuti (ovvero 300 s), perciò il numero di letture effettuate è pari a

$$\frac{300 \text{ s}}{200 \mu\text{s}} = 1.5 \times 10^6.$$

Ciascuno dei 1.5×10^6 segnali di trigger inviati ha attivato la lettura dei dati in un intervallo tem-

porale pari a 64 bc, ovvero $64 \cdot 25 \text{ ns} = 1.6 \text{ }\mu\text{s}$. Il tempo effettivo di acquisizione dati, pertanto, è pari a

$$1.5 \times 10^6 \cdot 1.6 \text{ }\mu\text{s} = 2.4 \text{ s} . \quad (6.2)$$

Il numero medio di elettroni atteso per ciascun pixel al termine dell'acquisizione è dato dal prodotto della (6.1) per la (6.2):

$$0.8 \text{ e}^- \cdot \text{s}^{-1} \cdot 2.4 \text{ s} = 1.9 \text{ e}^- .$$

Tale valore è inferiore rispetto agli 8 hit osservati in media nella mappa in figura 6.11 di un fattore 4. Ciò può essere spiegato considerando il fatto che gli elettroni, a causa della loro esigua massa, sono soggetti ad un elevato numero di eventi di scattering all'interno del Silicio; di conseguenza la loro interazione con il sensore non è localizzata. A titolo esemplificativo, in figura 6.12 è rappresentato il risultato dell'esposizione di un modulo costituito da un chip (Timepix3, [49]) con pixel di dimensioni $55 \text{ }\mu\text{m} \times 55 \text{ }\mu\text{m}$ ed un sensore con pixel di spessore $300 \text{ }\mu\text{m}$ ad una sorgente di ^{90}Sr ; è possibile notare come i singoli eventi (indicati con diversi colori) generino un segnale in più di un pixel. Tale risultato consente di giustificare da un punto

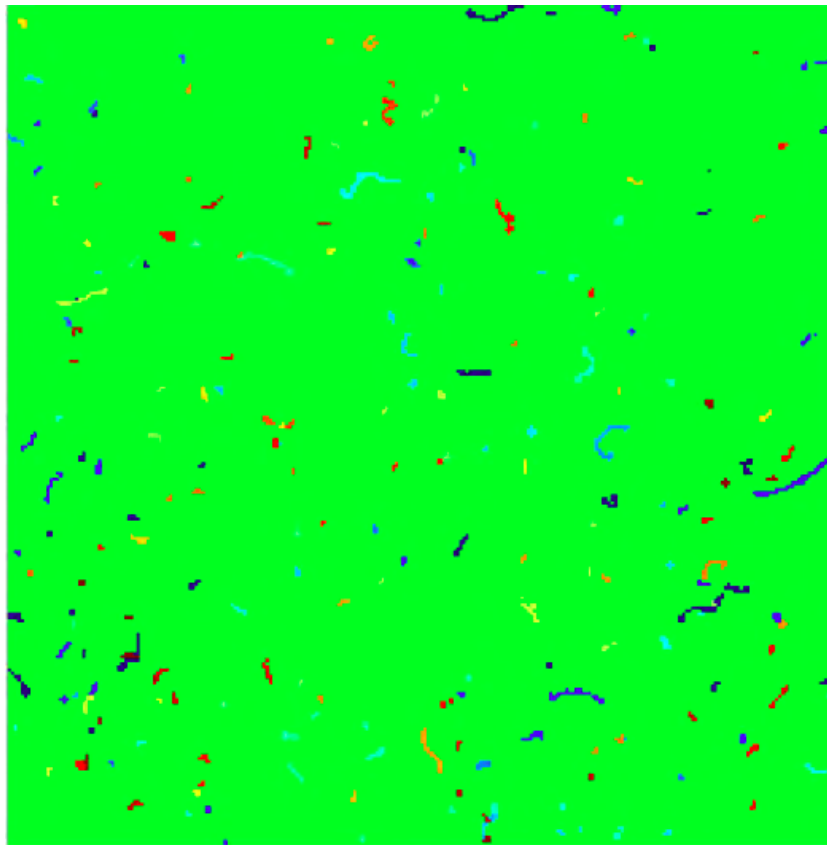


Figura 6.12: Mappa degli hit registrata da un modulo costituito da un chip (Timepix 3) con pixel di dimensioni $55 \text{ }\mu\text{m} \times 55 \text{ }\mu\text{m}$ ed un sensore con pixel di spessore $300 \text{ }\mu\text{m}$ esposto ad una sorgente di ^{90}Sr . Da [50].

di vista qualitativo la discrepanza osservata.

La distribuzione energetica degli elettroni emessi da una sorgente di ^{90}Sr è rappresentata in figura 6.13, mentre in figura 6.14 è rappresentata la perdita media di energia per ionizzazione per unità di lunghezza e di densità (espressa in $\text{MeV} \cdot \text{cm}^2 \cdot \text{g}^{-1}$ e spesso indicata come *stopping power*) subita da un elettrone che interagisce con il Silicio. Dal confronto tra le due figure si

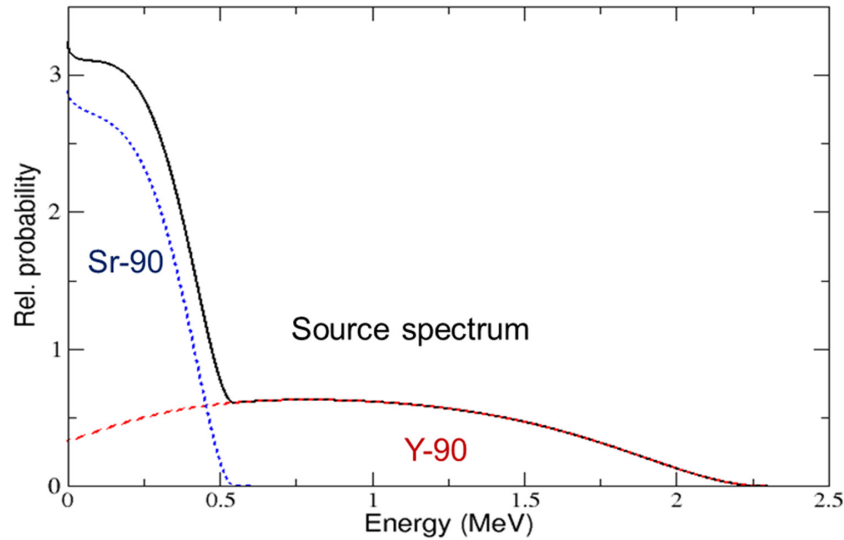


Figura 6.13: Distribuzione energetica degli elettroni emessi da una sorgente di ^{90}Sr . Da [51].

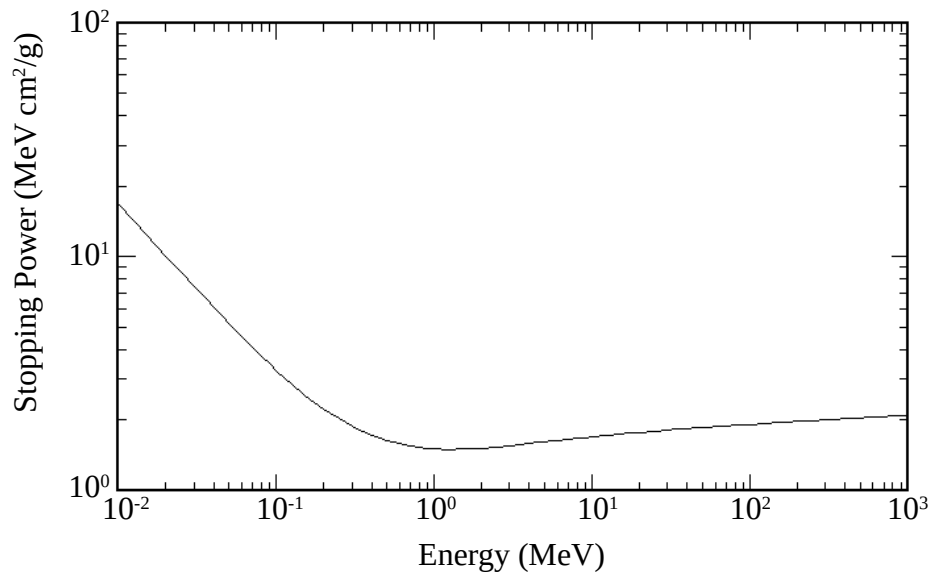


Figura 6.14: Illustrazione della perdita media di energia per ionizzazione per unità di lunghezza e di densità (*stopping power*) subita da un elettrone che interagisce con il Silicio in funzione della sua energia. Da [52].

osserva che tutti gli elettroni emessi dal decadimento dell' ^{90}Y attraversano il Silicio nelle condizioni di minima ionizzazione; un discorso analogo vale per gli elettroni prodotti dallo ^{90}Sr con energia maggiore di ~ 0.2 MeV. Si può quindi concludere che la maggior parte degli elettroni emessi dalla sorgente adoperata possono essere assimilati a particelle al minimo di ionizzazione.

Come introdotto nella sezione 5.5.5, il più probabile numero di elettroni eccitati in banda di conduzione al passaggio di una particella carica al minimo di ionizzazione in 150 μm di Silicio è pari a circa 10000. Tale affermazione può essere giustificata facendo riferimento al grafico illustrato in figura 6.15, che rappresenta la funzione di distribuzione del rapporto tra l'energia Δ persa da un pione di energia 500 MeV (che si trova nella condizione di minima ionizzazione) nell'interazione con il Silicio e lo spessore di materiale attraversato, x . La curva che descrive il

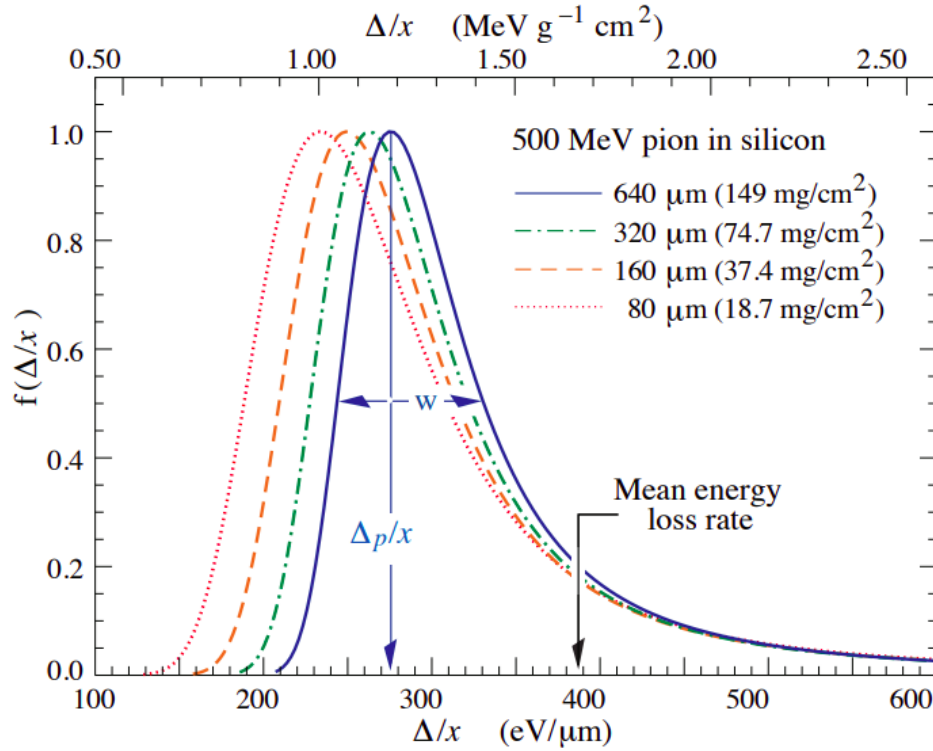


Figura 6.15: Distribuzione del rapporto tra la perdita di energia di un pione di energia 500 MeV che interagisce con il Silicio, Δ , e lo spessore di materiale attraversato, x . Da [53].

sensore in esame può essere ben approssimata con quella relativa ad uno spessore di 160 μm ; essa presenta un massimo in corrispondenza di 250 $\text{eV}/\mu\text{m}$, ovvero di un'energia di

$$250 \frac{\text{eV}}{\mu\text{m}} \cdot 150 \mu\text{m} = 37.5 \text{ keV}.$$

Considerando che l'energia media necessaria per creare una coppia elettrone - lacuna nel silicio è 3.6 eV, si ottiene che il più probabile numero di elettroni eccitati è

$$\frac{37500 \text{ eV}}{3.6 \text{ eV}} \sim 10417.$$

Dal momento che la procedura di tuning del CSA è stata eseguita in modo che una carica di 10000 e^- sia associata ad un ToT pari a 8 bc ci si può aspettare, in prima approssimazione, che i valori di ToT registrati da ciascun pixel siano descritti da una distribuzione analoga a quella rappresentata in figura 6.15 con media pari a 8 bc. Il risultato ottenuto al termine dell'acquisizione è rappresentato in figura 6.16. Il valor medio di questa distribuzione è significativamente

Distribuzione ToT

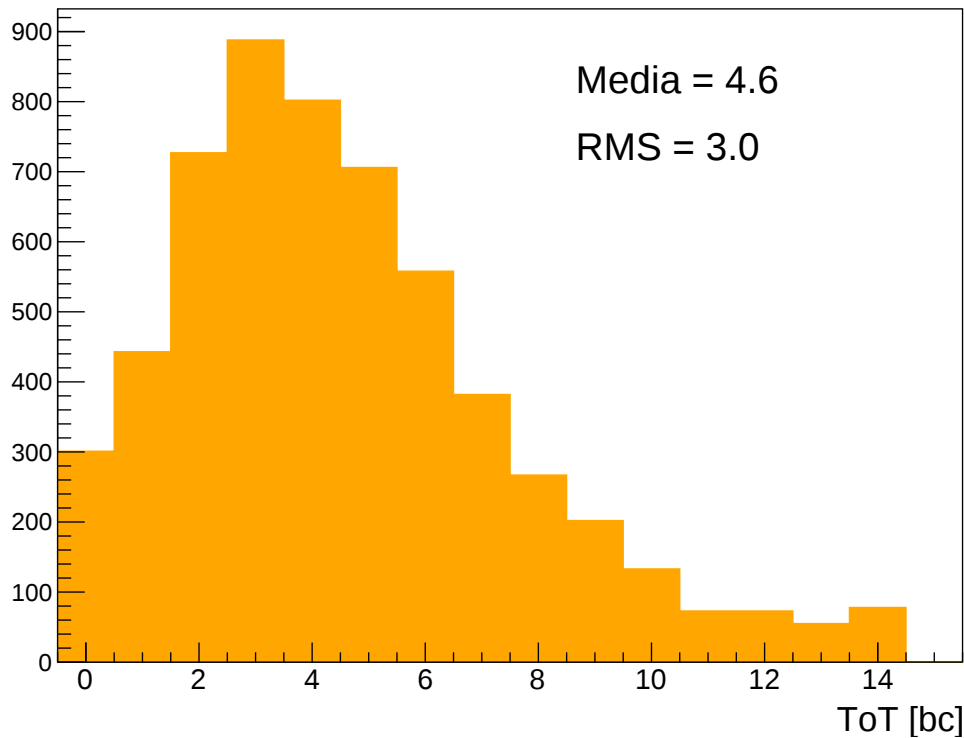


Figura 6.16: Distribuzione dei valori di ToT registrati da ciascun pixel durante l'esposizione del modulo ad una sorgente di ^{90}Sr .

inferiore rispetto a quello previsto. Ciò può essere spiegato considerando il fatto che gli elettroni, a causa dello scattering subito, si propagano nel sensore in diverse direzioni. Una possibile direzione di propagazione è quella parallela al lato più corto del pixel; in questo caso, la distanza percorsa dall'elettrone al suo interno è pari a $25\text{ }\mu\text{m}$, ovvero $1/6$ rispetto ai $150\text{ }\mu\text{m}$ associati ad un $\text{ToT} = 8\text{ bc}$, e pertanto il ToT del segnale generato è inferiore a 2 bc . La distribuzione di valor medio pari a $\text{ToT} = 8\text{ bc}$ attesa in prima approssimazione, quindi, deve essere corretta considerando effetti geometrici. Inoltre, fenomeni di *charge sharing*, termine con il quale si indica la diffusione degli elettroni generati in un pixel per effetto dell'eccitazione del mezzo in pixel contigui, comportano la degradazione del segnale, che è pertanto associato a bassi valori di ToT. La determinazione teorica della distribuzione attesa è un procedimento complesso che deve tener conto di diversi fattori, ed esula dallo scopo della presente trattazione.

Appendice A

Vincoli su λ dalle correzioni elettrodeboli al NLO al singolo Higgs

Il coefficiente di accoppiamento trilineare del bosone di Higgs interviene nelle correzioni elettrodeboli al NLO di processi di produzione e di decadimento di tale particella; in figura A.1 sono rappresentati alcuni dei diagrammi relativi a questi processi, nei quali l'accoppiamento trilineare è evidenziato in blu. Tali meccanismi possono dunque fornire informazioni riguardo al valore di

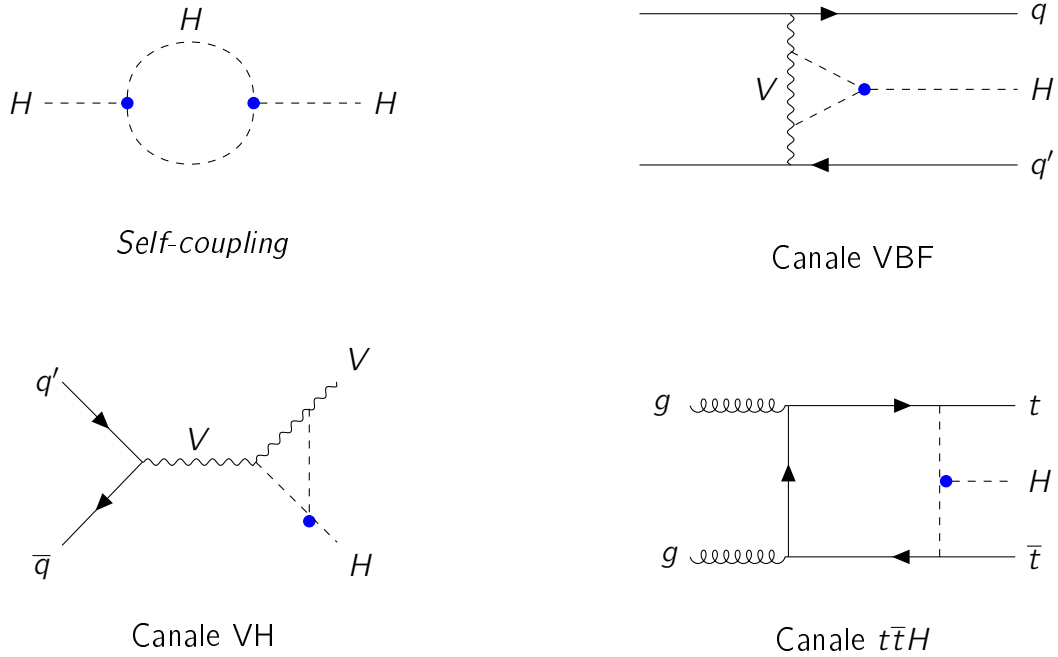


Figura A.1: Esempi di correzioni elettrodeboli al NLO dei processi di produzione di un singolo Higgs.

k_λ .

ATLAS ha svolto un'analisi volta a porre dei limiti sul valore di k_λ che si basa sull'idea che effetti di fisica BSM influenzino le sezioni d'urto differenziali ed il branching ratio dei vari processi che coinvolgono un singolo Higgs attraverso una modifica della costante di accoppiamento trilineare e delle altre costanti di accoppiamento dell'Higgs. L'analisi è stata effettuata sulla

base dei dati raccolti tra il 2015 ed il 2017, corrispondenti ad una luminosità integrata massima di 79.8 fb^{-1} (il suo valore è diverso a seconda del sottocampione considerato) [54].

È possibile definire un parametro che quantifichi la deviazione della sezione d'urto di produzione dell'Higgs in scenari di fisica BSM rispetto a quella prevista nel Modello Standard, indicato come *signal strength*. Questa assume, nel canale i , l'espressione

$$\mu_i(k_\lambda, k_i) = \frac{\sigma^{BSM}}{\sigma^{SM}} = Z_H^{BSM}(k_\lambda) \left[k_i^2 + \frac{(k_\lambda - 1)C_1^i}{K_{EW}^i} \right] \quad (\text{A.1})$$

dove $Z_H^{BSM}(k_\lambda)$ è definito come

$$Z_H^{BSM}(k_\lambda) = \frac{1}{1 - (k_\lambda^2 - 1)\delta Z_H} \quad \text{con} \quad \delta Z_H = -1.536 \times 10^{-3}.$$

Le correzioni elettrodeboli al NLO sono dunque il risultato di due contributi: uno universale (Z_H^{BSM}) ed uno dipendente dal processo e dalla cinematica in esame (il termine tra parentesi quadre nella (A.1)). Il valore k_i indica l'entità delle correzioni elettrodeboli al LO alla sezione d'urto del processo di produzione in canale i nell'ipotesi $k_\lambda = 1$ ($k_i^2 = \frac{\sigma_{LO,i}^{BSM}}{\sigma_{LO,i}^{SM}}$, con $k_\lambda = 1$), K_{EW}^i quella al NLO ($K_{EW}^i = \frac{\sigma_{NLO,i}^{SM}}{\sigma_{LO,i}^{SM}}$) e C_1^i è la componente dipendente dal processo e dalla cinematica. Nell'analisi condotta da ATLAS sono stati considerati unicamente i fattori correttivi alle costanti di accoppiamento dell'Higgs ai fermioni (k_F) ed ai bosoni vettori massivi (k_V). Pertanto, il fattore k_i^2 assume per ogni canale di produzione considerato le espressioni indicate in tabella A.1.

Canale di produzione	ggF	VBF	VH	$t\bar{t}H$
k_i^2	k_F^2	k_V^2	k_V^2	k_F^2

Tabella A.1: Espressione del fattore k_i^2 per i canali di produzione dell'Higgs considerati.

Analogamente alle sezioni d'urto di produzione, la signal strength relativa al processo di decadimento dell'Higgs nel canale f , $\mu_f(k_\lambda, k_f) = \frac{BR_f^{BSM}}{BR_f^{SM}}$, acquista una dipendenza da k_λ e da $k_f^2 = \frac{BR_{LO,f}^{BSM}}{BR_{LO,f}^{SM}}$. Quest'ultimo assume in ciascun canale di decadimento considerato le espressioni indicate in tabella A.2.

Canale	$H \rightarrow \gamma\gamma$	$H \rightarrow WW^*$	$H \rightarrow ZZ^*$	$H \rightarrow b\bar{b}$	$H \rightarrow \tau\tau$
k_f^2	$1.59k_V^2 + 0.07k_F^2 - 0.67k_Vk_F$	k_V^2	k_V^2	k_F^2	k_F^2

Tabella A.2: Espressione del fattore k_f^2 per i canali di decadimento dell'Higgs considerati.

La dipendenza delle sezioni d'urto totali e dei branching ratios dal valore di k_λ ottenuta è rappresentata in figura A.2, rispettivamente a sinistra e a destra. Le curve sono normalizzate alle espressioni assunte nel Modello Standard e sono suddivise in base ai canali di produzione e decadimento considerati.

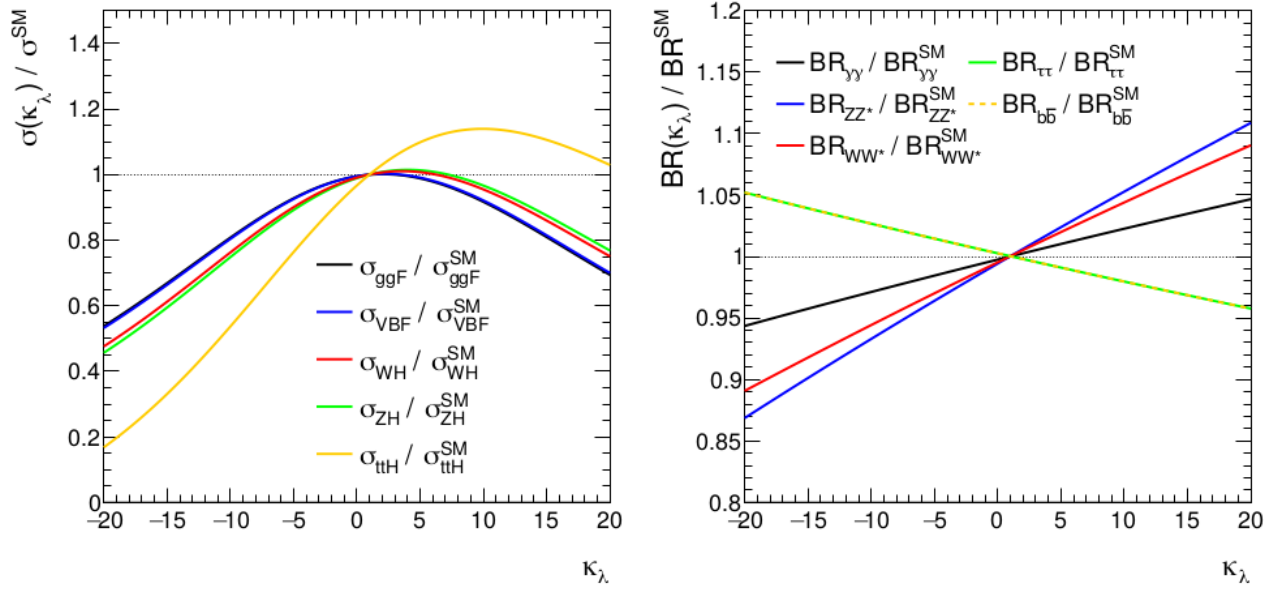


Figura A.2: Dipendenza delle sezioni d'urto (sinistra) e dei branching ratios (destra) dei processi di produzione e decadimento dell'Higgs dal valore di k_λ . Da [54].

Il parametro utilizzato per individuare possibili deviazioni dal Modello Standard è

$$\mu_{if}(k_\lambda, k_F, k_V) = \mu_i(k_\lambda, k_F, k_V) \times \mu_f(k_\lambda, k_F, k_V).$$

È possibile effettuare un fit di verosimiglianza combinato alle sezioni d'urto di produzione di Higgs nei diversi modi indicati in tabella A.1 sfruttando i canali di decadimento più abbondanti indicati in tabella A.2 per porre dei vincoli all'intervallo di variabilità di k_λ assumendo $k_F = k_V = 1$. Il risultato ottenuto è rappresentato in figura A.3. Il valore di *best fit* è

$$k_\lambda = 4.0^{+4.3}_{-4.1} = 4.0^{+3.7}_{-3.6}(\text{stat.})^{+1.6}_{-1.5}(\text{exp.})^{+1.3}_{-0.9}(\text{sig. th.})^{+0.8}_{-0.9}(\text{bkg. th.}).$$

L'incertezza totale è stata decomposta nella componente statistica (*stat.*, dominante), in quella sperimentale (*exp*) ed in quella teorica proveniente dalla modellizzazione del segnale (*sig. th.*) e del fondo (*bkg. th.*).

L'intervallo di variabilità ammesso per k_λ al livello di confidenza del 95% sulla base dei dati osservati è:

$$-3.2 < k_\lambda < 11.9.$$

Tale risultato è compatibile con quello ottenuto ricercando la produzione non risonante di coppie di bosoni di Higgs (indicato nell'espressione (1.2)) e dimostra come lo studio delle correzioni elettrodeboli al NLO in processi che coinvolgono un solo Higgs possa fornire risultati di significatività comparabile con il primo tipo di analisi.

Per testare scenari BSM in cui gli effetti di nuova fisica influenzino l'accoppiamento dell'Higgs ai fermioni oppure ai bosoni vettori massivi oltre che l'accoppiamento trilineare è possibile effettuare un fit volto a determinare i possibili intervalli di variabilità di k_λ e k_F una volta fissato

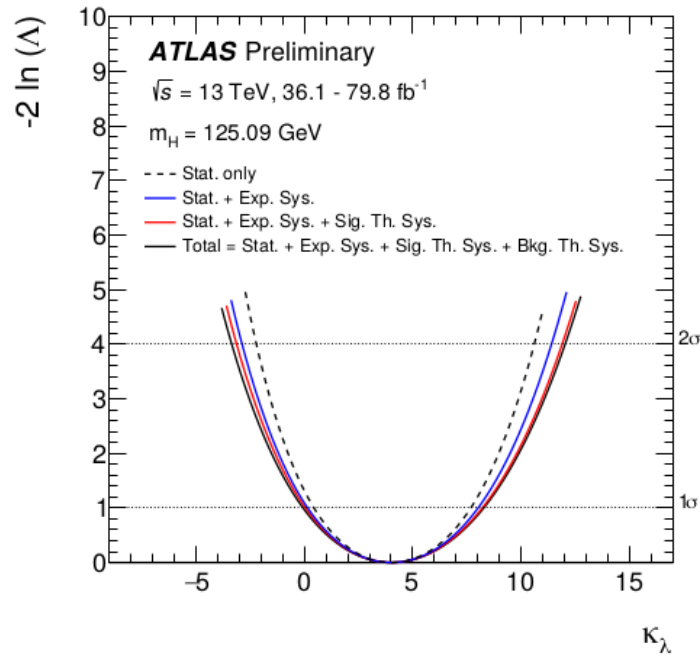


Figura A.3: Fit di verosimiglianza in funzione di k_λ effettuato considerando diverse fonti di errore. Da [54].

$k_V = 1$ ed uno in cui il parametro variabile è k_V , mentre $k_F = 1$. Lo spazio dei parametri compatibile con i dati al termine delle procedure di fit è rappresentato in figura A.4, in cui sono stati indicati i valori attesi nel Modello Standard, il risultato del best fit e le regioni consentite al livello di confidenza del 68% e del 95%. Si osserva che l'aggiunta di un grado di libertà al fit

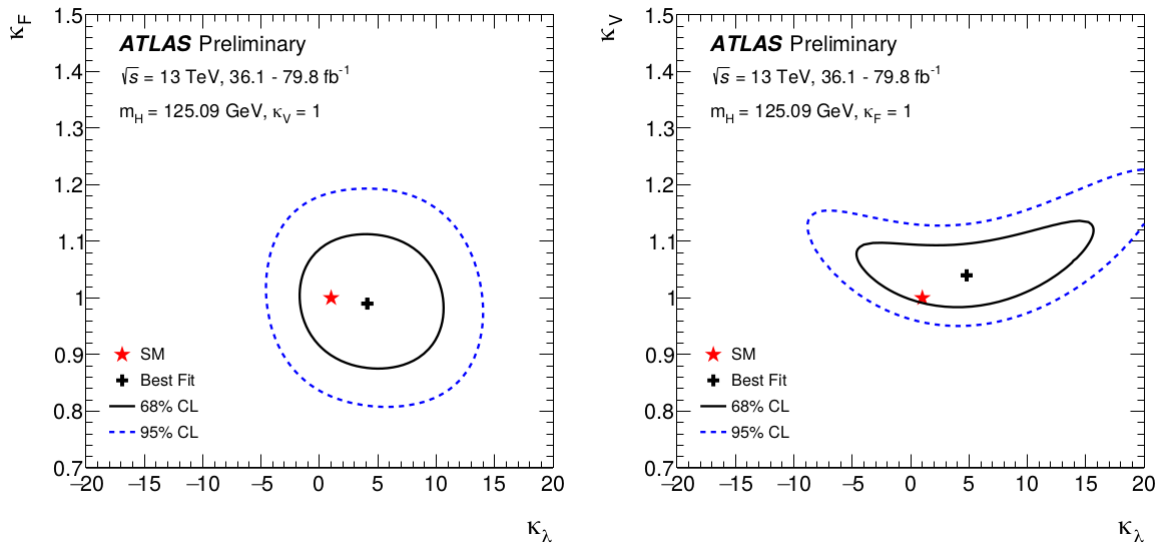


Figura A.4: Valori ammissibili dei parametri k_λ e k_F (sinistra) e k_λ e k_V (destra) ottenuti al termine della procedura di fit. Da [54].

precedente rende i risultati ottenuti in termini di k_λ meno restrittivi, in particolare nel caso in

cui si ammette la variabilità di k_V ; ciò rende tale analisi meno competitiva rispetto alla ricerca di coppie di Higgs.

Appendice B

Il rivelatore ATLAS

Con la sua lunghezza di 46 m ed il suo diametro di 25 m, ATLAS è il più grande dei quattro rivelatori di particelle operativi ad LHC; situato 100 m sotto il livello del suolo, ha una forma cilindrica ed un peso di circa 7000 tonnellate. Una vista tridimensionale del rivelatore è rappresentata in figura B.1.

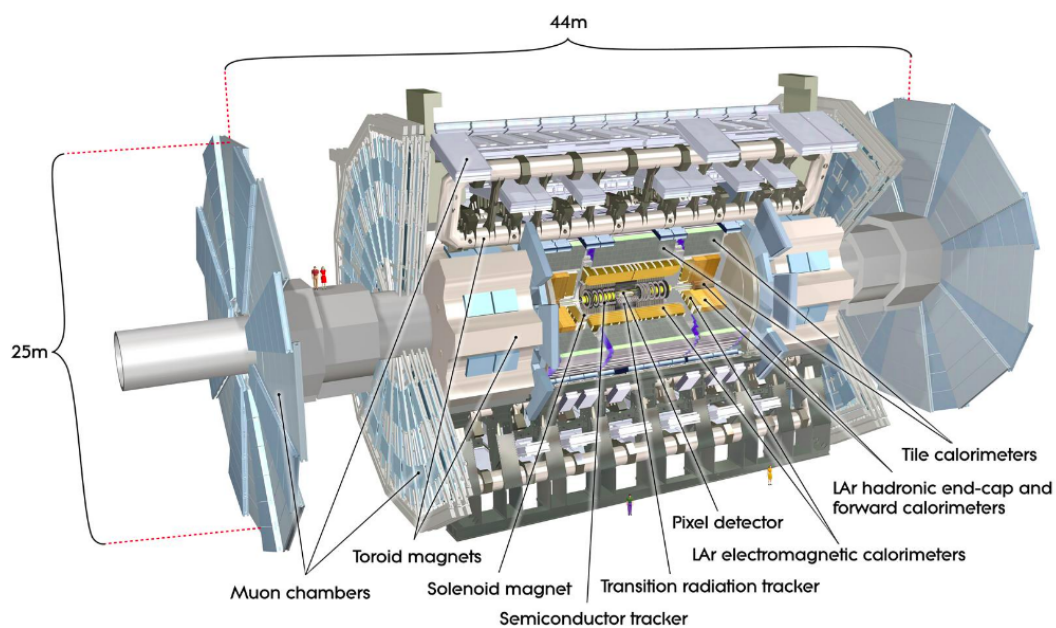


Figura B.1: Vista tridimensionale del rivelatore ATLAS.

La struttura di ATLAS rispecchia quella caratteristica di un rivelatore di particelle: è costituito da una successione di diversi sottorivelatori concentrici, organizzati in modo tale che ognuno risponda ad una specifica esigenza di misura. Procedendo dall'interno verso l'esterno si susseguono:

- l'Inner Detector, utilizzato per misurare direzione, momento e carica delle particelle cariche prodotte dall'interazione protone - protone. È costituito da un sottorivelatore a pixel di Silicio, uno a strip di Silicio ed un rivelatore di radiazione di transizione;
- un solenoide che genera un campo magnetico in cui è immerso l'Inner Detector;

- un calorimetro elettromagnetico ed uno adronico, entrambi a campionamento, utilizzati per misurare l'energia delle particelle tramite il loro assorbimento completo;
- lo spettrometro per muoni, utilizzato per misurare il momento di tali particelle. È collocato nella parte più esterna dal momento che i muoni sono le particelle più penetranti (ad esclusione dei neutrini, che comunque non sono rivelati direttamente a causa della loro bassissima sezione d'urto e sono identificati tramite il mancato bilanciamento dell'energia nel piano trasverso);
- un campo magnetico toroidale, da cui il nome del rivelatore, che curva le traiettorie dei muoni permettendo la misura del loro momento.

Appendice C

Rivelatori a pixel monolitici

Negli ultimi decenni si è svolta un'intensa attività di ricerca volta a sviluppare una tecnologia che consenta di realizzare rivelatori a pixel che superino i limiti di quelli ibridi, descritti nella sezione 3.5.1.

In questo tipo di rivelatore, detto *monolitico*, l'elettronica di front-end è integrata all'interno del sensore; ciò consente di ridurre il rumore elettronico, grazie all'eliminazione dei bump-bond, il numero di moduli danneggiati a causa di imperfezioni nella realizzazione dei bump-bond e lo spessore del rivelatore. Tuttavia, questa tecnologia è meno resistente, allo stato attuale, rispetto a quella ibrida ai danni indotti dall'interazione con elevate dosi di radiazioni. La risoluzione di questo problema rappresenta il motivo principale per cui numerose collaborazioni svolgono attività di ricerca e sviluppo sui rivelatori monolitici.

Due possibili geometrie di rivelatori monolitici sono rappresentate in figura C.1. Nel primo

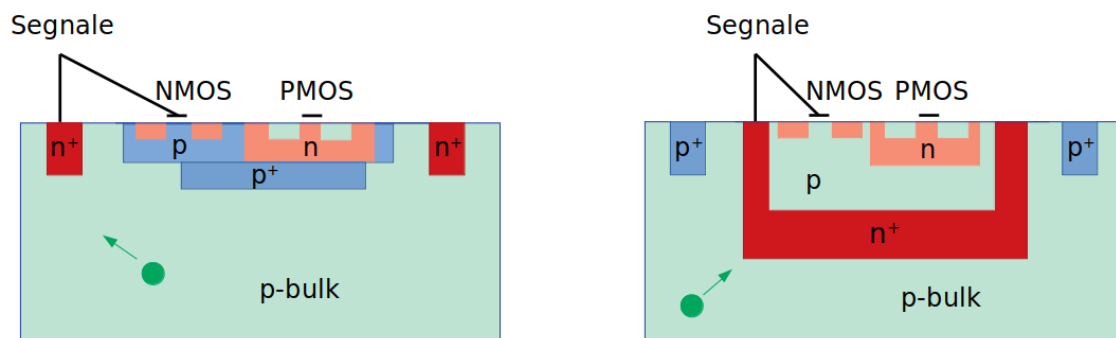


Figura C.1: Possibili varianti di rivelatori monolitici, in cui gli elettrodi di raccolta della carica sono rappresentati dagli impianti di tipo n^+ . A sinistra l'elettronica di front-end è integrata tra due elettrodi adiacenti, a destra al loro interno.

caso i pixel sono costituiti da impianti di tipo n^+ su un p -bulk e rappresentano dunque gli elettrodi di raccolta della carica. Ciò permette di estendere la regione di svuotamento soprattutto nel bulk, come indicato dalle equazioni (3.7). L'elettronica di front-end è integrata tra due elettrodi adiacenti; la tecnologia adoperata per la sua realizzazione è quella CMOS, basata sull'utilizzo di transistor MOSFET a canale N (NMOS) e P (PMOS). Ognuno di essi è opportunamente schermato dalla regione di svuotamento tramite strati di drogaggio opposto a quello del loro bulk (impianti di tipo p per i transistor NMOS e di tipo n per i PMOS). I tratti neri sulla superficie

superiore del bulk rappresentano i rivestimenti metallici. Le giunzioni pn hanno superficie molto ridotta e ciò consente di ottenere prestazioni veloci ed un basso rumore elettronico, come conseguenza della (3.3) e della (3.8). Tuttavia, ciò implica anche che in alcune condizioni, ad esempio in seguito ai danni indotti dalla radiazione, non si riesca a svuotare l'intero bulk. Questo problema è risolto nella seconda tipologia di rivelatori monolitici, rappresentata a destra in figura C.1. In questo caso su un substrato di tipo p è realizzato un impianto n^+ , il quale funge da elettrodo di raccolta di carica. I due impianti p^+ rappresentano due pixel adiacenti. L'elettronica di front-end è realizzata utilizzando la stessa tecnologia CMOS, ma è collocata all'interno dell'elettrodo. La superficie della giunzione pn è molto maggiore rispetto a quella della geometria precedente e questo consente di svuotare l'intero bulk. Tuttavia, ciò comporta anche un aumento della capacità e pertanto un rallentamento del dispositivo e l'aumento del rumore elettronico.

La ricerca di soluzioni ai problemi descritti ha portato allo sviluppo di diverse varianti di rivelatori monolitici. Una di queste è stata selezionata dall'esperimento ALICE [55], ad LHC, per il miglioramento del suo sistema di tracciamento [56] (*Inner Tracking System*, ITS), iniziato al termine del 2018.

Conclusioni

L'attività svolta per la stesura di questa tesi consiste nella caratterizzazione di un prototipo dei chip di lettura che saranno utilizzati nel rivelatore a pixel di ATLAS durante la fase di presa dati denominata High-Luminosity LHC (HL-LHC), che prenderà avvio nel 2027. Esso rappresenta il rivelatore interno del nuovo tracciatore in Silicio, detto Inner Tracker (ITk), che sostituirà quello attualmente utilizzato dall'esperimento, inadeguato a sostenere l'aumento di luminosità previsto in questa fase.

HL-LHC è ad oggi l'unico programma di ricerca in fisica delle particelle agli acceleratori dedicato all'indagine alla frontiera dell'energia che sia approvato e finanziato. La presa dati si estenderà per oltre 10 anni e vedrà impegnata una comunità internazionale estremamente ampia che si sviluppa a partire da quella attualmente coinvolta negli esperimenti ATLAS e CMS (oltre 5000 fisici, ingegneri e tecnici appartenenti a centinaia di istituzioni di ricerca). Il programma scientifico di HL-LHC, oltre a contemplare la generale ricerca di nuovi fenomeni, prevede la realizzazione di misure di precisione del Modello Standard e, tra queste, la misura dell'intensità dell'interazione fra tre bosoni di Higgs, parametro fondamentale del potenziale del campo di Higgs. Tale obiettivo può essere raggiunto misurando la sezione d'urto di produzione di coppie di bosoni di Higgs, un processo estremamente raro secondo le predizioni del Modello Standard che, tuttavia, ad HL-LHC potrà essere osservato con abbondanza sufficiente da consentire una misura combinata dei due esperimenti che raggiunga la significatività di 5σ .

Dal momento che tale obiettivo ha motivato l'aumento di luminosità e ha dettato la progettazione dei rivelatori degli esperimenti, questo elaborato ha avuto inizio con una discussione dei risultati ottenuti fino ad ora dalla collaborazione ATLAS nella ricerca di produzione di coppie di Higgs con i dati raccolti all'energia nel centro di massa di 13 TeV. Tali ricerche hanno fornito esito negativo, in accordo con quanto previsto dal Modello Standard, ed hanno prodotto dei limiti di esclusione per il valore del coefficiente di accoppiamento trilineare dell'Higgs.

La sensibilità delle analisi già pubblicate e di quelle che in futuro saranno condotte sui dati di HL-LHC dipende fortemente dall'efficacia con cui si riesce ad affrontare una problematica sperimentale cruciale negli esperimenti ai collisori adronici: il riconoscimento di jet originati da quark b (b -tagging) e la loro discriminazione rispetto ai jet prodotti da quark più leggeri o gluoni che affollano gli eventi prodotti nelle collisioni multiple tra i fasci di protoni accelerati ad LHC. L'importanza del b -tagging nella fisica dell'Higgs dipende dal fatto che il modo di decadimento di gran lunga più abbondante di questo bosone (quasi 60% del totale) è quello nello stato finale $b\bar{b}$. Pertanto, in questa tesi è stata discussa l'evoluzione degli algoritmi di b -tagging nel corso degli ultimi anni e l'impatto sulle loro prestazioni determinato dall'evoluzione della struttura del rivelatore di tracciamento, dall'introduzione dell'Insertable B-Layer prima della fase di presa dati iniziata nel 2015 fino al passaggio dall'attuale rivelatore all'ITk per HL-LHC. Sulla base dell'efficienza di tracciamento e di b -tagging attese con l'ITk sui dati che ATLAS acquisirà durante HL-LHC è possibile quindi discutere la prospettiva di una misura del

coefficiente di accoppiamento trilineare dell'Higgs con la statistica finale della fase di presa dati ad alta luminosità. Tali tematiche, affrontate nella prima parte di questo lavoro, consentono di apprezzare l'importanza di progettare l'Inner Tracker in modo che abbia delle caratteristiche che permettano di migliorare l'efficienza di b -tagging, aumentando dunque la possibilità di individuare coppie di bosoni di Higgs.

L'incremento di luminosità previsto impone la progettazione di nuovi rivelatori a pixel, più adatti a lavorare in ambienti caratterizzati da un'alta frequenza di produzione di particelle e da un'elevata dose di radiazioni ionizzanti rispetto a quelli attualmente installati nel rivelatore ATLAS. Le modifiche previste coinvolgono entrambe le componenti dei moduli ibridi da realizzare: il sensore, nel quale avviene la generazione del segnale, ed il chip di front-end, che provvede alla digitalizzazione di tale segnale ed alla sua trasmissione al mondo esterno. L'attività illustrata in questa tesi si focalizza su un prototipo di tali chip di front-end, detto RD53A; essa è stata svolta presso il Laboratorio INFN "Costruzione grandi apparati" del Dipartimento di Matematica e Fisica dell'Università del Salento, utilizzando un sistema di acquisizione dati sviluppato all'interno della collaborazione ATLAS per la caratterizzazione dei chip RD53A. Diversi campioni dei prototipi sono stati distribuiti ai gruppi di ricerca impegnati nel progetto del rivelatore a pixel di ITk in modo che potessero svolgere test con l'obiettivo di verificarne la funzionalità e di individuare la migliore logica di front-end tra le tre implementate.

Durante questo lavoro il sistema di acquisizione dati è stato adattato al setup allestito in laboratorio, il quale presenta delle specificità rispetto a quello utilizzato da altri gruppi. In questo modo è stato possibile eseguire sul chip a disposizione la procedura di calibrazione elaborata dagli sviluppatori del software di acquisizione dati, che consiste nell'analizzare la risposta del chip all'iniezione di una quantità nota di carica all'ingresso del front-end e nel modificare i suoi parametri di configurazione in modo da ottenere la risposta desiderata. I risultati ottenuti sono stati confrontati con quelli ricavati precedentemente da altri gruppi e hanno contribuito a confermare la loro validità. Inoltre, sono state effettuate delle misure non previste da tale procedura di calibrazione con l'obiettivo di verificare con maggiore dettaglio la funzionalità dell'elettronica di front-end.

Sulla base delle misure illustrate in questo lavoro di tesi è possibile concludere che la componente digitale dei circuiti integrati nel chip analizzato funziona in maniera corretta; la componente analogica, invece, manifesta dei problemi in alcuni pixel. Questi sono contenuti nella regione della matrice del chip nella quale è implementata la logica di front-end di tipo differenziale; la loro presenza era prevista sulla base dei risultati ottenuti precedentemente da altri gruppi e l'analisi dei segnali registrati dal chip è stata condotta escludendo quelli provenienti da tali pixel. Nella stessa regione del chip è stato osservato anche che la media della distribuzione dei valori di *Time over Threshold*, termine con il quale si indica la durata dell'intervallo in cui l'ampiezza del segnale registrato da un pixel è superiore alla soglia impostata dall'utente, non raggiunge mai il massimo valore possibile. Anche tale comportamento è stato riscontrato nei risultati ottenuti da altri gruppi ed è pertanto da intendersi come una conferma di un fenomeno noto. Inoltre, è stata verificata la linearità dei tre convertitori digitale-analogico che forniscono a ciascuno dei front-end implementati la soglia con la quale confrontare l'ampiezza del segnale rivelato da ogni pixel per discriminare quello generato da una particella ionizzante da quello dovuto al rumore. Infine, è stato indagato il comportamento del sensore integrato sul chip prima di esporlo ad una sorgente radioattiva. I risultati di tale analisi hanno indicato il danneggiamento della maggior parte dei bump-bond utilizzati per il collegamento del sensore al chip e l'incompatibilità dell'andamento della corrente nel sensore al variare della tensione applicata con quello atteso.

Questi comportamenti anomali erano previsti, poiché era noto che il rivelatore a disposizione per le misure è gravemente compromesso per quel che riguarda la funzionalità del sensore e la sua connessione al chip. L'analisi qualitativa dei risultati ottenuti esponendo il modulo ad una sorgente di ^{90}Sr è risultata compatibile con le previsioni teoriche.

Le misure condotte per la stesura di questa tesi hanno contribuito allo sviluppo del setup che il gruppo INFN di Lecce utilizzerà per l'assemblaggio di una sezione del rivelatore a pixel che sarà integrato nell'Inner Tracker di ATLAS e all'acquisizione dell'abilità di caratterizzare sistematicamente tali moduli gestendo un sistema di acquisizione complesso ed estremamente innovativo.

Bibliografia

- [1] ATLAS Collaboration, *The ATLAS Experiment at the CERN Large Hadron Collider*, [JINST **3** \(2008\) S08003](#).
- [2] CMS Collaboration, *The CMS experiment at the CERN LHC*, [JINST **3** \(2008\) S08004](#).
- [3] L. Evans and P. Bryant, *LHC Machine*, [JINST **3** \(2008\) S08001](#).
- [4] F. Englert and R. Brout, *Broken symmetry and the mass of gauge vector mesons*, [Phys. Rev. Letters **13** \(1964\) 321](#).
- [5] P. Higgs, *Broken symmetries and the masses of gauge bosons* [Phys. Rev. Lett. **13** \(1964\) 508](#).
- [6] ATLAS Collaboration, *Observation of a New Particle in the Search for the Standard Model Higgs Boson with the ATLAS Detector at the LHC*, [Phys. Lett. B **716** \(2012\) 1-29](#).
- [7] CMS Collaboration, *Observation of a new boson at a mass of 125 GeV with the CMS experiment at the LHC* [Phys. Lett. B **716** \(2012\) 30-61](#).
- [8] M. Tanabashi et al. (Particle Data Group), *Status of Higgs Boson Physics*, [Phys. Rev. D **98**, 030001 \(2018\) and 2019 update](#).
- [9] ATLAS Collaboration, *Study of the spin of the new boson with up to 25 fb⁻¹ of ATLAS data*, [ATLAS-CONF-2013-040 \(2013\)](#).
- [10] D. de Florian et al. (LHC Higgs Cross Section Working Group), *Handbook of LHC Higgs cross sections: 4. Deciphering the nature of the Higgs sector*, arXiv: [1610.07922v2 \[hep-ph\] \(2017\)](#).
- [11] R. Frederix, S. Frixione, V. Hirschi, F. Maltoni, O. Mattelaer, P. Torrielli, E. Vryonidou and M. Zaro, *Higgs pair production at the LHC with NLO and parton-shower effects*, arXiv: [1401.7340v2 \[hep-ph\] \(2014\)](#).
- [12] ATLAS Collaboration, *Search for pair production of Higgs bosons in the $b\bar{b}b\bar{b}$ final state using proton-proton collisions at $\sqrt{s} = 13$ TeV with the ATLAS detector*, [JHEP **01** \(2019\) 030](#).
- [13] ATLAS Collaboration, *Search for Resonant and Nonresonant Higgs Boson Pair Production in the $b\bar{b}\tau^+\tau^-$ Decay Channel in pp Collisions at $\sqrt{s} = 13$ TeV with the ATLAS detector*, [Phys. Rev. Lett. **121** \(2018\) 191801](#).

-
- [14] ATLAS Collaboration, *Search for non-resonant Higgs boson pair production in the $\gamma\gamma b\bar{b}$ final state with 13 TeV pp collision data collected by the ATLAS experiment*, [JHEP **11** \(2018\) 040](#).
- [15] ATLAS Collaboration, *Search for Higgs boson pair production in the $b\bar{b}WW^*$ decay mode at $\sqrt{s} = 13\text{ TeV}$ with the ATLAS detector*, [JHEP **04** \(2019\) 092](#).
- [16] ATLAS Collaboration, *Search for non-resonant Higgs boson pair production in the $WW^{(*)}WW^{(*)}$ decay channel using ATLAS data recorded at $\sqrt{s} = 13\text{ TeV}$* , [JHEP **05** \(2019\) 124](#).
- [17] ATLAS Collaboration, *Search for Higgs boson pair production in the $\gamma\gamma WW^*$ channel using pp collision data recorded at $\sqrt{s} = 13\text{ TeV}$ with the ATLAS detector*, [Eur. Phys. J. C **78** \(2018\) 1007](#).
- [18] ATLAS Collaboration, *Search for the $HH \rightarrow b\bar{b}b\bar{b}$ process via vector-boson fusion production using proton–proton collisions at $\sqrt{s} = 13\text{ TeV}$ with the ATLAS detector*, [arXiv: 2001.05178 \[hep-ex\] \(2020\)](#).
- [19] ATLAS Collaboration, *Combination of searches for Higgs boson pairs in pp collisions at $\sqrt{s} = 13\text{ TeV}$ with the ATLAS detector*, [Phys. Lett. B **800** \(2020\) 135103](#).
- [20] CMS Collaboration, *Combination of Searches for Higgs Boson Pair Production in Proton-Proton Collisions at $\sqrt{s} = 13\text{ TeV}$* , [Phys. Rev. Lett. **122** \(2019\) 121803](#).
- [21] P. A. Zyla et al. (Particle Data Group), to be published in [Prog. Theor. Exp. Phys. **2020**, 083C01 \(2020\)](#).
- [22] C. Patrignani et al. (Particle Data Group), *Review of Particle Physics*, [Chin. Phys. C, **2016**, 40\(10\): 100001 \(2016\)](#).
- [23] ATLAS Collaboration, *Tagging and suppression of pileup jets with the ATLAS detector*, [ATLAS-CONF-2014-018 \(2018\)](#).
- [24] S. Grancagnolo on behalf of the ATLAS and CMS Collaborations, *Top quark pair-production cross section measurements at the LHC*, [ATL-PHYS-PROC-2019-057 \(2019\)](#).
- [25] ATLAS Collaboration, *Optimisation and performance studies of the ATLAS b-tagging algorithms for the 2017-18 LHC run*, [ATL-PHYS-PUB-2017-013 \(2017\)](#).
- [26] ATLAS Collaboration, *Expected performance of the ATLAS b-tagging algorithms in Run-2*, [ATL-PHYS-PUB-2015-022 \(2015\)](#).
- [27] ATLAS Collaboration, *Secondary vertex finding for jet flavour identification with the ATLAS detector*, [ATL-PHYS-PUB-2017-011 \(2017\)](#).
- [28] ATLAS Collaboration, *Topological b-hadron decay reconstruction and identification of b-jets with the JetFitter package in the ATLAS experiment at the LHC*, [ATL-PHYS-PUB-2018-025 \(2018\)](#).

-
- [29] Frühwirth, R., *Application of Kalman filtering to track and vertex fitting*, [Nucl. Instrum. Meth. A](#) **262** 444 (1987).
- [30] ATLAS IBL Collaboration, *Production and Integration of the ATLAS Insertable B-Layer*, arXiv: [1803.00844v3 \[physics.ins-det\]](#) (2018).
- [31] *Impact Parameter Resolution*, [IDTR-2015-007](#) (2015).
- [32] *Plots of b-tagging performance before and after the installation of the Insertable B-Layer*, [FTAG-2017-005](#) (2017).
- [33] *ITk Pixel Layout Updates*, [ITK-2020-002](#) (2020).
- [34] ATLAS Collaboration, *Expected Tracking Performance of the ATLAS Inner Tracker at the HL-LHC*, [ATL-PHYS-PUB-2019-014](#) (2019).
- [35] ATLAS Collaboration, *Technical Design Report for the ATLAS Inner Tracker Pixel Detector*, [ATLAS-TDR-030, 2017](#) (2017).
- [36] ATLAS Collaboration, *Measurement prospects of the pair production and self-coupling of the Higgs boson with the ATLAS experiment at the HL-LHC*, [ATL-PHYS-PUB-2018-053](#) (2018).
- [37] ATLAS Collaboration, *ATLAS b-jet identification performance and efficiency measurement with $t\bar{t}$ events in pp collisions at $\sqrt{s} = 13$ TeV*, arXiv: [1907.05120v2 \[hep-ex\]](#) (2020).
- [38] M. Cooke, *Monitoring the radiation damage of the ATLAS pixel detector*, [Nucl. Instrum. Meth. A](#) **718** 299-301 (2013).
- [39] W. Shockley, *Currents to Conductors Induced by a Moving Point Charge*, [Journal of Applied Physics](#) **9**, 635 (1938).
- [40] S. Ramo, *Currents Induced by Electron Motion*, [Proceedings of the IRE](#) **27** 584 (1939).
- [41] *RD-53 Collaboration Home Page*, <http://rd53.web.cern.ch/>.
- [42] M. Garcia-Sciveres, *The RD53A Integrated Circuit*, [CERN-RD53-PUB-17-001](#) (2019).
- [43] F. Krummenacher, *Pixel detectors with local intelligence: an IC designer point of view*, [Nucl. Instr. Meth. A](#) **305**, pp. 527-532 (1991).
- [44] Xilinx, *Aurora 64B/66B Protocol Specification*, [SP011 \(v1.3\) October 1](#) (2014).
- [45] RD53 Collaboration, *RD53B users guide: Introduction to RD53B pixel chip architecture, features and recommendations for use in pixel detector systems (Version 0.3)* [CERN-RD53-PUB-19-002](#) (2020).
- [46] N. Giangiacomi, et al., *General purpose readout board π LUP: overview and results*, arXiv: [1806.08858v1 \[physics.ins-det\]](#) (2018).
- [47] T. Heim, *Performance of the Insertable B-Layer for the ATLAS Pixel Detector during Quality Assurance and a Novel Pixel Detector Readout Concept based on PCIe*, [University of Wuppertal, Dissertation](#) (2015).

-
- [48] ATLAS Collaboration, *Technical Design Report for the ATLAS Inner Tracker Pixel Detector*, [ATLAS-TDR-030 \(2018\)](#).
- [49] T. Poikela et al., *Timepix3: a 65K channel hybrid pixel readout chip with simultaneous ToA/ToT and sparse readout*, [JINST **9** \(2014\) C05013](#).
- [50] M. Campbell et al., *Charged particle detection using the timepix and timepix3 chips and future plans*, https://www2.physics.ox.ac.uk/sites/default/files/2012-03-27/mccampbell_oxford_pdf_14057.pdf.
- [51] T. Roy, F. Tessier, M. McEwen, *A system for the measurement of electron stopping powers: proof of principle using a pure β -emitting source*, [Radiation Physics and Chemistry **149** \(2018\) 134-141](#).
- [52] ESTAR (stopping-power and range tables for electrons), <https://physics.nist.gov/PhysRefData/Star/Text/ESTAR.html>.
- [53] P.A. Zyla et al. (Particle Data Group), *Passage of Particles Through Matter* [Prog. Theor. Exp. Phys. **2020**, 083C01 \(2020\)](#).
- [54] ATLAS Collaboration, *Constraint of the Higgs boson self-coupling from Higgs boson differential production and decay measurements*, [ATL-PHYS-PUB-2019-009 \(2019\)](#).
- [55] ALICE Collaboration, *The ALICE experiment at the CERN LHC*, [JINST **3** S08002 \(2008\)](#).
- [56] ALICE Collaboration, *Technical Design Report for the Upgrade of the ALICE Inner Tracking System*, [J. Phys. G: Nucl. Part. Phys. **41** 087002 \(2014\)](#).