



Entwicklung eines Triggersystems zum Testen und Kalibrieren der Supermodule des ALICE-TRD

Diplomarbeit in Physik

vorgelegt von

Jonas Anielski

Institut für Kernphysik
der Westfälischen Wilhelms-Universität
Münster

Februar 2011

Inhaltsverzeichnis

1. Einleitung	1
2. Grundlagen	3
2.1. Das Standardmodell	4
2.2. Quark-Gluon Plasma	6
2.3. Detektion geladener Teilchen	11
3. Der LHC und ALICE	17
3.1. Large Hadron Collider	17
3.2. ALICE	19
3.3. Das ALICE-Triggersystem	22
4. Transition Radiation Detector	31
4.1. Motivation für den TRD	31
4.2. Übergangsstrahlung	32
4.3. Aufbau des TRD	32
4.4. Funktionsprinzip	35
4.5. Triggersystem beim TRD	37
4.6. Detektorelektronik	40
5. Kosmische Strahlung	47
5.1. Primäre kosmische Strahlung	47
5.2. Sekundäre kosmische Strahlung	47
6. Detektion kosmischer Teilchen in Münster	51
6.1. Motivation für einen Trigger auf kosmische Teilchen	51
6.2. Aufbau zur Detektion kosmischer Teilchen	52
6.3. Koinzidenzbedingung	58
6.4. Diskriminatorschwellen der Trefi-Module	60
7. Triggersystem	71
7.1. Motivation für eine neue Triggereinheit	71
7.2. Aufbau in der Montagehalle in Münster	73

7.3. Zustandsmaschine	77
7.4. Interner Logikanalysator	81
8. Trigger	87
8.1. Trigger für kosmische Teilchen	87
8.2. Testtrigger	98
9. Benutzerschnittstelle	109
9.1. Kommunikation mit dem FPGA	109
9.2. Aufbau der Benutzerschnittstelle	110
10. Zusammenfassung und Ausblick	119
 Anhang	 123
A. Tabellen zum Triggersystem bei ALICE	123
B. Logikpegel	125
C. Szintillatoren und Trefi-Module	127
D. Schaltpläne der NIM-LVDS- und LVDS-NIM-Wandler	129
E. Register der Triggereinheit	135
 Danksagung	 149

1. Einleitung

Mit *A Large Ion Collider Experiment* (ALICE) am *Large Hadron Collider* (LHC) des CERN¹ soll mit Hilfe von Bleikollisionen ein *Quark–Gluon Plasma* (QGP) erzeugt und untersucht werden. Der ALICE–Detektor besteht aus Subdetektoren, die für verschiedene Aufgaben wie die Teilchenidentifizierung oder Impulsmessung der entstehenden Teilchen optimiert sind. Da die Subdetektoren verschieden schnell ausgelesen werden können und die Datenmenge reduziert werden muss, wurde für ALICE ein mehrstufiges Triggersystem entwickelt, mit dem physikalisch interessante Ereignisse ausgewählt werden.

Einer der Subdetektoren ist der Übergangsstrahlungsdetektor (TRD: kurz für *Transition Radiation Detector*), der Elektronen mit hohen Impulsen identifizieren soll und einen schnellen Beitrag zum Triggersystem von ALICE liefert. Der TRD ist in 18 sogenannte *Supermodule* unterteilt, die einzeln in Münster montiert und getestet werden.

Im Rahmen dieser Arbeit wurde ein Triggersystem für Münster entwickelt, mit dem die Supermodule kalibriert und getestet werden können. Dazu wurde der Aufbau aus [Bat07] teilweise übernommen und weiterentwickelt. Priorität bei der Entwicklung des neuen Systems hatte die Stabilität bei der Datennahme und den Tests, da die Stromversorgung der Detektorelektronik durch falsche Trigger des alten Aufbaus oft zusammengebrochen ist.

Zusätzlich sollte der Funktionsumfang erweitert und die Benutzerfreundlichkeit erhöht werden, indem sich das Triggersystem vollständig über ein TCP/IP–Netzwerk steuern lässt.

Die zentrale Triggereinheit des neuen Systems stellt ein *V1495 General Purpose Board* von Caen mit einem frei programmierbaren Altera Cyclone FPGA dar. Für diesen wurde ein Design in *Very High Speed Integrated Circuit Hardware Description Language* (VHDL) entwickelt und getestet.

Nach einer kurzen historischen Einführung wird in Kapitel 2 auf die physikalischen Grundlagen von ALICE und auf die Detektion geladener Teilchen eingegangen. Anschließend werden der LHC und ALICE in Kapitel 3 erläutert. Dort werden die Subdetektoren kurz vorgestellt und ein Überblick über das globale Triggersystem von ALICE gegeben. Der TRD wird in Kapitel 4 detailliert besprochen. Hier wird noch einmal das Triggersystem speziell für den TRD aufgegriffen, da sich das Triggersystem in Münster stark daran orientiert. In Kapitel 5 wird die natürliche kosmische Strahlung diskutiert, da sie zur Kalibrierung der Supermodule verwendet wird.

Der Nachweis kosmischer Teilchen zur Triggererzeugung in Münster wird in Kapitel 6 erörtert. In

¹ European Organization for Nuclear Research. CERN stand ursprünglich für Conseil Européen pour la Recherche Nucléaire.

Kapitel 7 wird dann die zentrale Triggereinheit, basierend auf dem V1495, vorgestellt. Außerdem wird ein interner Logikanalysator vorgestellt, der zur Analyse und Fehlersuche bei den Triggern bzw. der Triggereinheit dient. Die einzelnen Trigger, ihre Funktion und Performance werden in Kapitel 8 beschrieben. In Kapitel 9 wird dann die Benutzerschnittstelle, mit der alle Funktionen der Triggereinheit überwacht und kontrolliert werden können, vorgestellt. Zum Schluss wird in Kapitel 10 eine Zusammenfassung dieser Arbeit und ein Ausblick gegeben.

2. Grundlagen

Nach einer kurzen historischen Einleitung wird in diesem Kapitel das Standardmodell der Elementarteilchenphysik vorgestellt. Im Anschluss wird das sogenannte *Quark-Gluon Plasma* (QGP), ein Phasenzustand von Materie, in dem sich das Universum in den ersten Mikrosekunden nach dem Urknall befand, beschrieben und diskutiert, wie solch ein Zustand im Labor erzeugt und nachgewiesen werden kann. Daraufhin wird noch der Energieverlust von geladenen Teilchen in Materie diskutiert und es werden zwei Detektortypen vorgestellt.

Schon seit langer Zeit suchen Wissenschaftler nach den Bausteinen der Materie. Dabei wurden im Laufe der Zeit immer kleinere Konstituenten gefunden. Ende des 19. Jahrhunderts war bekannt, dass Materie aus Atomen besteht, und die verschiedenen Atome wurden in das Periodensystem der Elemente eingeordnet. Auf Grund der periodischen Eigenschaften wurde angenommen, dass Atome sich aus kleineren Bausteinen zusammensetzen.

1897 entdeckte Thomson das Elektron, indem er Kathodenstrahlen untersuchte. Rutherford fand bald darauf bei Streuversuchen heraus, dass sich fast die gesamte Masse von Atomen in einem kleinen, schweren, positiv geladenen Kern befindet, der von Elektronen umgeben ist. 1919 entdeckte er, dass der Kern von Wasserstoff auch in anderen Elementen enthalten ist und identifizierte damit das zweite Elementarteilchen: das Proton.

Nachdem Pauli 1930 das Neutrino postulierte, um das Energiespektrum und die vermeintliche Verletzung der Drehimpulserhaltung beim β -Zerfall zu erklären, dauerte es noch 26 Jahre es experimentell nachzuweisen, da es nur schwach wechselwirkt. In der Zwischenzeit (1932) wurde bereits das Neutron durch Chadwick entdeckt.

In den 1950ern wurde in Beschleunigerexperimenten eine Vielzahl von weiteren Teilchen gefunden. In diesem sogenannten *Teilchenzoo* befinden sich neben anderen Teilchen viele Hadronen. Das sind Teilchen, die der starken Wechselwirkung unterliegen. Gell-Man und Zweig schlugen vor, dass es sich bei den Hadronen nicht um Elementarteilchen handelt, sondern dass sie sich aus Partonen zusammensetzen. Durch tiefinelastische Streuung wurde dies experimentell am *Stanford Linear Accelerator* (SLAC) bestätigt. Heute ist bekannt, dass sich Hadronen aus *Quarks* und *Gluonen*¹ zusammensetzen. Innerhalb der Hadronen wird zwischen *Mesonen*, die aus einem Quark und einem Antiquark, und den *Baryonen*, die aus drei Quarks bzw. drei Antiquarks bestehen, unterschieden.

Quarks sind die kleinsten Bausteine der Materie die bisher bekannt sind.

¹ Wird von Quarks und Gluonen zusammen gesprochen, werden sie aus historischen Gründen auch Partonen genannt.

2.1. Das Standardmodell

Die bekannten Elementarteilchen und ihre Wechselwirkungen werden im Standardmodell der Elementarteilchenphysik zusammengefasst. Es beschreibt die *Eichbosonen*, die die fundamentalen Kräfte vermitteln und drei Generationen von *Fermionen*², denen jeweils zwei Quarks und zwei Leptonen zugeordnet sind.

2.1.1. Quarks und Leptonen

Es existieren 6 Quarks: *Up*-, *Down*-, *Charm*-, *Strange*-, *Top*- und *Bottom*-Quark (u, d, c, s, t, b). Diese tragen neben der elektrischen Ladung von $+2/3 e$ bzw. $-1/3 e$ noch eine Farbladung *rot*, *grün* oder *blau*. Sie können nicht einzeln beobachtet werden, was als *Confinement* bezeichnet wird (vgl. Kap. 2.2).

Des Weiteren existieren 6 Leptonen. Im Einzelnen sind das: Elektron, Myon, Tauon und die zugeordneten Neutrinos: Elektron-, Myon- und Tauon-Neutrino.

Zu jedem Teilchen gibt es zusätzlich das entsprechende Antiteilchen, das die gleiche Masse, aber die entgegengesetzte elektrische Ladung, Farbladung und dritte Komponente des schwachen Isospins hat. In Tabelle 2.1 ist eine Auflistung der 12 Fermionen und ihrer Eigenschaften zu sehen.

	Familie	Name	Ladung(e)	I_3	Wechselwirkung	Masse (MeV/c^2)
Quarks	1	u	$+2/3$	$+1/2$	EM, schwach, stark	$1,7 - 3,3$
		d	$-1/3$	$-1/2$	EM, schwach, stark	$4,1 - 5,8$
	2	c	$+2/3$	$+1/2$	EM, schwach, stark	$(1,18 - 1,34) \cdot 10^3$
		s	$-1/3$	$-1/2$	EM, schwach, stark	$72 - 122$
	3	t	$+2/3$	$+1/2$	EM, schwach, stark	$(168,8 - 174,2) \cdot 10^3$
		b	$-1/3$	$-1/2$	EM, schwach, stark	$(4,13 - 4,85) \cdot 10^3$
Leptonen	1	e	-1	$-1/2$	EM, schwach	0,511
		ν_e	0	$+1/2$	schwach	$< 2 \cdot 10^{-6}$
	2	μ	-1	$-1/2$	EM, schwach	105
		ν_μ	0	$+1/2$	schwach	$< 0,17 \cdot 10^{-6}$
	3	τ	-1	$-1/2$	EM, schwach	$1,78 \cdot 10^3$
		ν_τ	0	$+1/2$	schwach	$< 18,2$

Tab. 2.1.: Quarks und Leptonen im Standardmodell [Pov06, PDG10].

2.1.2. Wechselwirkungen und Eichbosonen

Es gibt vier fundamentale Wechselwirkungen. Das sind die elektromagnetische, schwache und starke Wechselwirkung sowie die Gravitation. Sie sind in Tabelle 2.2 aufgeführt. Im Standardmodell

² Fermionen sind Teilchen, die einen halbzahligen Spin haben. In einem Ensemble aus Fermionen müssen sie sich in mindestens einer Quantenzahl unterscheiden. Bosonen, die einen ganzzahligen Spin haben, unterliegen dieser Beschränkung nicht.

sind nur die ersten drei enthalten, da die Gravitation auf subatomarer Ebene vernachlässigt werden kann. Für zwei Protonen ist sie um einen Faktor von ca. 10^{-36} [WW06] schwächer als die elektromagnetische Wechselwirkung.

Elektromagnetische Wechselwirkung

Die elektromagnetische Wechselwirkung wird durch Photonen vermittelt. Diese sind masselos und koppeln an die elektrische Ladung von Teilchen, tragen aber selbst keine elektrische Ladung. Dadurch hat die elektromagnetische Wechselwirkung eine unendliche Reichweite. Ihr Verhalten wird durch die Maxwell-Gleichungen beschrieben.

Schwache Wechselwirkung

Die schwache Wechselwirkung wird durch $W^\pm$ - und Z^0 -Bosonen vermittelt, die an die schwache Ladung³ koppeln. Die Austauschbosonen sind mit ca. 80 bzw. 90 MeV relativ schwer und können auf Grund der Heisenbergschen Unschärferelation nur für kurze Zeit existieren. Daher ist die Reichweite der schwachen Wechselwirkung auf weniger als 1 fm beschränkt.

Starke Wechselwirkung

Die starke Wechselwirkung wird durch die Quantenchromodynamik (QCD) beschrieben. Sie wird durch 8 Gluonen, die an die Farbladung der Quarks koppeln, vermittelt. Ein Gluon setzt sich immer aus einer Farb- und einer Antifarbladung zusammen. Da sie eine Farbladung tragen, können Gluonen auch untereinander wechselwirken. Das Potential V der starken Wechselwirkung in einem Meson in Abhängigkeit des Abstandes r zwischen Quark und Antiquark (siehe Abb. 2.1 b)) wird beschrieben durch [Ber02]:

$$V(r) = -\frac{4}{3} \frac{\alpha_s}{r} + \sigma r \quad (2.1)$$

$$\alpha_s \propto \frac{1}{\ln(Q^2)} \quad (2.2)$$

Hier ist $\alpha_s(Q^2)$ die Kopplungskonstante (siehe Abb. 2.1 a)) der starken Wechselwirkung, $\sqrt{Q^2}$ der Impulsübertrag und σ ein konstanter Faktor. Die Kopplungskonstante wird klein, wenn der Impulsübertrag groß ist. Da höhere Impulsüberträge einen geringeren Abstand implizieren, sind Quarks bei sehr geringen Abständen quasi-frei. Diese Eigenschaft wird *asymptotische Freiheit* genannt. Bei großen Abständen dominiert der lineare Term im Potential und das Potential wird sehr groß. Quarks können unter normalen Umständen nicht einzeln beobachtet werden, da die benötigte Energie, um zwei Quarks voneinander zu trennen, mit dem Abstand wächst. Wird versucht zwei

³ Die schwache Ladung entspricht der dritten Komponente des Isospins I_3

2. Grundlagen

Quarks voneinander zu trennen, wird ein neues Quark–Antiquark–Paar erzeugt, sobald die Energie ausreicht. Aus dem ursprünglichen Quark–Antiquark–Paar werden dann zwei Paare. Es existieren also nur farbneutrale Hadronen. Dieses Phänomen wird *Confinement* genannt.

Wechselwirkung	Reichweite	Eichboson	Ladung (e)	Masse (GeV/c^2)
Starke WW	≈ 1 fm	8 Gluonen	0	0
Schwache WW	$\ll 1$ fm	$W^\pm$	± 1	80,22
		Z^0	0	91,19
Elektr.-mag. WW	∞	Photon	0	0
Gravitation	∞	Graviton	0	0

Tab. 2.2.: Wechselwirkungen im Standardmodell und die Gravitation [WW06, TM04].

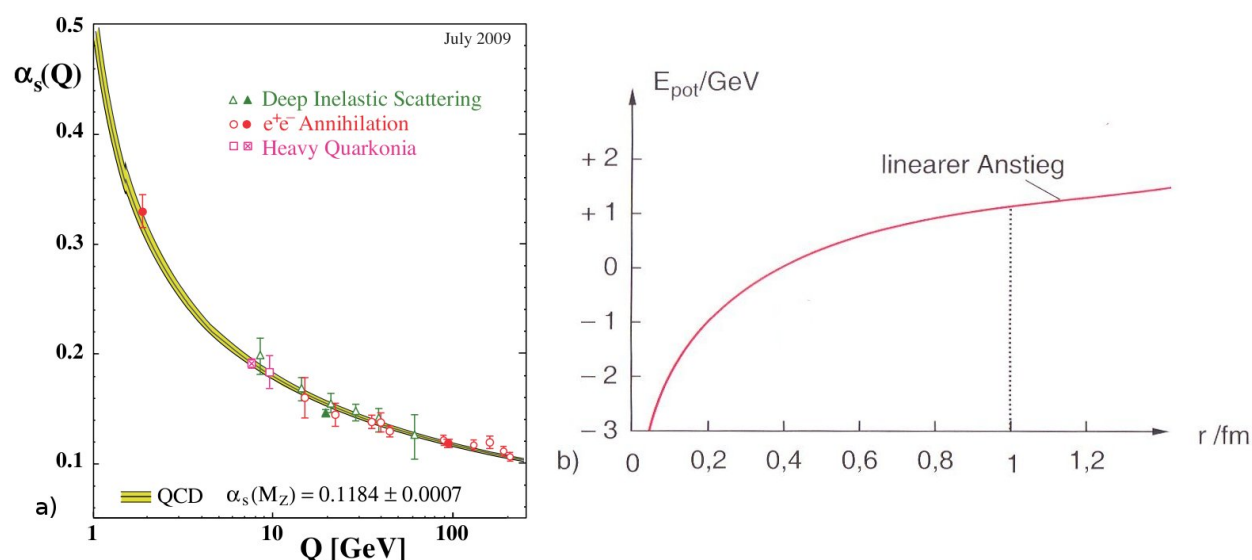


Abb. 2.1.: a) Die Kopplungskonstante α_s der starken Wechselwirkung in Abhängigkeit vom Impulsübertrag (logarithmisch aufgetragen). Für große Impulsüberträge wird die Kopplungskonstante klein. Bild aus [PDG10] b) Das Potential eines Quark–Antiquark–Paares in Abhängigkeit ihres Abstandes. Im Gegensatz zu den anderen Wechselwirkungen steigt das Potential mit steigendem Abstand. Ab $r = 1$ dominiert der lineare Term. Bild aus [Dem10].

2.2. Quark-Gluon Plasma

Wird baryonische Materie soweit erhitzt oder komprimiert, dass sich die Wellenfunktionen der einzelnen Baryonen überlagern, findet ein Phasenübergang statt, so dass Quarks und Gluonen ungebunden sind. Diese Phase wird Quark-Gluon-Plasma (QGP) genannt.

Mit der Gitter-QCD kann dieser Phasenübergang theoretisch beschrieben werden. Der neue Materiezustand lässt sich durch thermodynamische Größen wie die Temperatur und Dichte klassifizieren.

Im Wesentlichen gibt es zwei Parameter, die die Phase bestimmen [YHM05]:

- Dichte: Wird die Dichte auf ein Mehrfaches der Dichte von Kernmaterie erhöht, so dass sich die Kerne überlagern, entsteht ein kaltes QGP.
- Temperatur: Wird die Temperatur soweit erhöht, dass eine kritische Temperatur T_c überschritten wird, wird ein heißes QGP erzeugt.

In Abb. 2.2 ist ein Phasendiagramm, in dem die Temperatur T gegen das chemische Potential μ_B aufgetragen ist, gezeigt. Das chemische Potential ist eine Maßeinheit für die Baryonen-Dichte. Für ein Plasma, bestehend aus zwei leichten Quarks (up, down) und einem schweren (strange), wird eine Phasenübergangstemperatur T_c von 173 MeV mit einem systematischen Fehler von 10% bei einem chemischen Potential $\mu_B = 0$ erwartet [BMS07].

Kühlt ein QGP ab, bzw. dehnt sich aus, findet die Hadronisierung statt: aus den Quarks und Gluonen bilden sich Hadronen. Der Übergang bei niedrigem chemischen Potential und hoher Temperatur wird als kontinuierlich angenommen und entspricht dem Phasenübergang, der nach dem Urknall stattgefunden hat und bei ALICE beobachtet werden soll. Bei niedrigeren Temperaturen findet ein Phasenübergang 1.Ordnung statt. Für große chemische Potentiale und kleine Temperaturen wird eine farbsupraleitende Phase vermutet (vgl. Abb. 2.2).

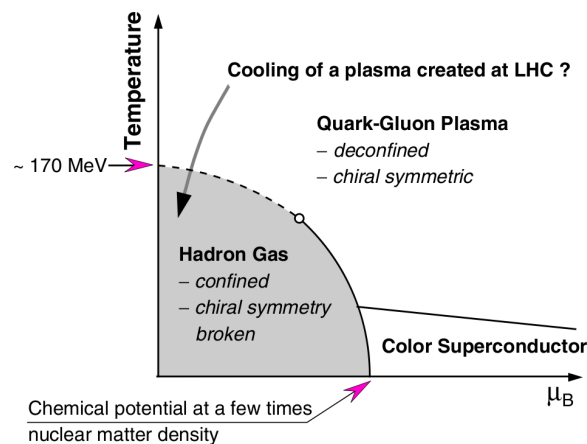


Abb. 2.2.: Phasendiagramm für den Übergang zum QGP. Bei einem chemischen Potential $\mu_B = 0$ wird eine Übergangstemperatur von ca. 170 MeV erwartet. Die durchgezogene Linie bedeutet einen Phasenübergang 1.Ordnung, während die gestrichelte Linie einen kontinuierlichen Übergang anzeigt. [ALI04].

2.2.1. Erzeugung eines QGP

Es wird angenommen, dass das Universum für kurze Zeit nach dem Urknall als QGP existiert hat. Um die Entstehung des Universums weiter zu erforschen, soll die QGP-Phase im Labor weiter

untersucht werden. Es ist jedoch schwierig, die benötigten Energiedichten im Labor zu erzeugen. Die Lösung dazu ist, dass schwere Atomkerne in Teilchenbeschleunigern auf nahezu Lichtgeschwindigkeit beschleunigt und aufeinander geschossen werden. Treffen sie aufeinander, erhitzt sich die Kernmaterie stark.

In Abb. 2.3 ist eine simulierte Kollision zweier Blei-Kerne zu sehen. Auf Grund ihrer hohen Geschwindigkeit sind sie stark Lorentz-kontrahiert und wirken daher wie Scheiben, die aufeinander zufliegen. Wenn die Ionen aufeinander treffen, finden harte Parton-Parton-Stöße statt. Es werden einige Tausend oder Zehntausend [BMS07] Gluonen und Quarks erzeugt, die sich bei der extrem hohen Energiedichte frei bewegen. Nach ca. 1 fm c^{-1} findet sich ein lokales thermisches Gleichgewicht ein. Die große Anzahl an Teilchen und das thermische Gleichgewicht sind notwendig, damit für das QGP thermodynamische Eigenschaften wie Phasenübergänge oder die Temperatur betrachtet werden können [BMS07].

Das QGP dehnt sich dann schnell aus und kühlt ab. Fällt die Temperatur unter die kritische Temperatur T_c , findet ein Phasenübergang statt: aus den Quarks und Gluonen bilden sich wieder Hadronen. Dieser Vorgang wird Hadronisierung genannt.

Der Zeitpunkt ab dem keine inelastischen Stöße mehr unter den Hadronen stattfinden, wird *chemisches Ausfrieren* genannt. Das bedeutet, dass die Zusammensetzung der Hadronen sich nicht mehr ändert. Ist die mittlere freie Weglänge der Teilchen größer als die Ausdehnung des Systems, entweichen sie der Teilchenwolke. Dieser Zeitpunkt wird *thermisches Ausfrieren* genannt. Ab diesem Moment sind die Impulse der Hadronen fest. Entstandene Teilchen können von den umliegenden Detektoren erfasst und identifiziert werden.

Der entscheidende Faktor, ob bei einer Kollision die benötigten Teilchendichten bzw. Temperaturen für ein QGP erreicht werden können, ist die Energie, auf die die Beschleuniger die Ionen bringen⁴. Im Laufe der Zeit wurden immer bessere Beschleuniger gebaut. So erreicht der LHC, der in Kapitel 3 diskutiert wird, Energien von $\sqrt{s_{NN}} = 5,5 \text{ TeV}$ bei Kollisionen von Bleikernen, was die Energie des bis dahin besten Beschleunigers um das ca. 30-fache überragt. $\sqrt{s_{NN}}$ ist die Energie pro Nukleonen-Paar im Schwerpunktsystem, womit auch im Folgenden die Energien von Kollisionen bzw. Beschleunigern beschrieben werden.

Mit Schwerionenkollisionen konnten in der Vergangenheit bereits eindeutige Hinweise für die Erzeugung eines QGP im Labor gesammelt werden. Einige der charakteristischen Hinweise werden im Folgenden vorgestellt.

2.2.2. Signaturen des QGP

Ein QGP kann nicht direkt nachgewiesen werden, sondern nur indirekt über Teilchen, die bei der Kollision entstehen und durch die umliegenden Detektoren identifiziert werden. Charakteristische Unterschiede in den Observablen bei Kollisionen, bei denen ein QGP entsteht und solchen, bei denen

⁴ Es gibt auch andere wichtige Faktoren. Insbesondere die Anzahl der Ionen, die im Beschleuniger gespeichert sind und wie gut sie an der Kollisionsstelle kollimiert werden, ist entscheidend für die Statistik, die erreicht wird.

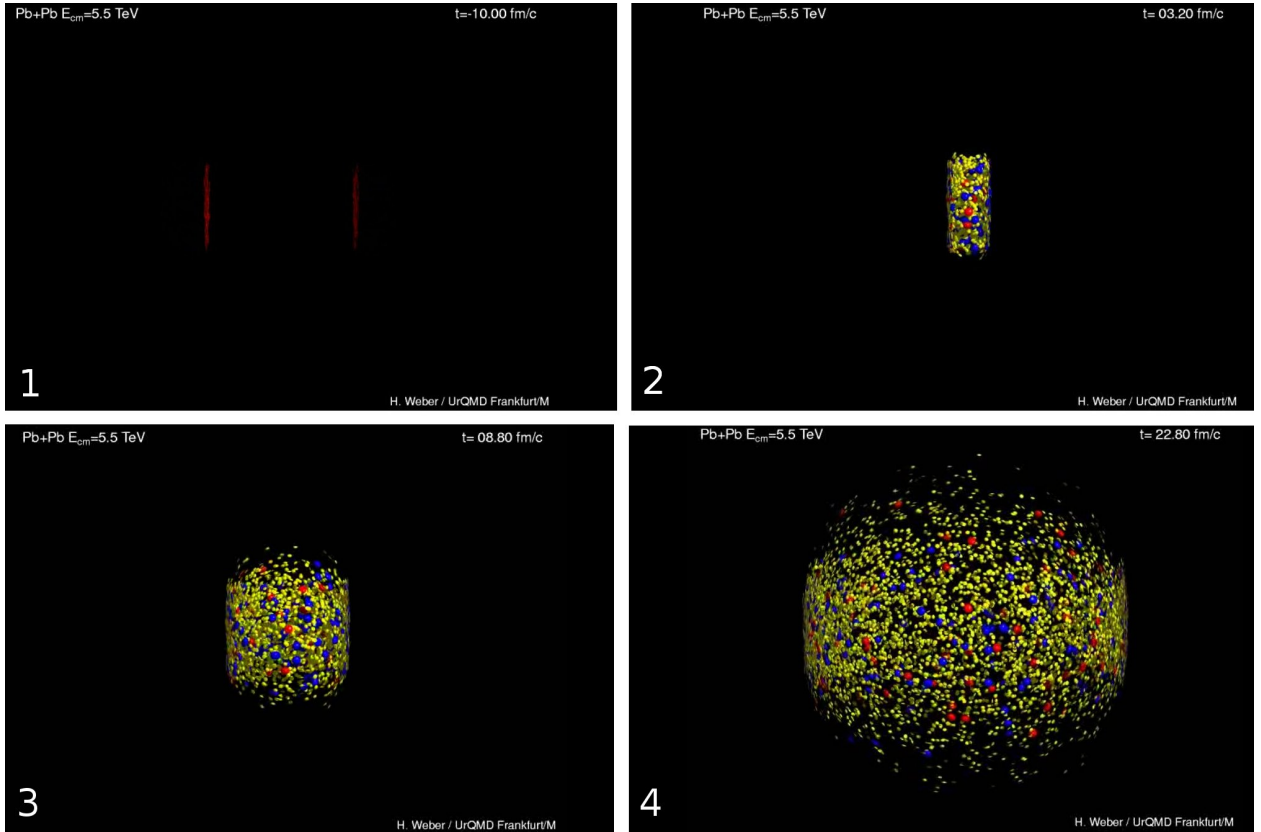


Abb. 2.3.: Simulation einer Kollision zweier Lorentz-kontrahierter Bleikerne bei $\sqrt{s_{NN}} = 5,5$ TeV. Die Nukleonen der Kerne sind hier in rot, Mesonen in gelb und angeregte Baryonen in blau dargestellt. Bild aus [Web10]

kein QGP erzeugt wird, werden Signaturen genannt.

Mit vorherigen Beschleunigern wie dem *Super Proton Synchrotron* (SPS) am CERN und dem *Relativistic Heavy Ion Collider* (RHIC) am *Brookhaven National Laboratory* (BNL) wurden bereits einige Signaturen, die typisch für ein QGP sind, nachgewiesen.

Es wird erwartet, dass die Eigenschaften des QGP mit dem LHC genauer als bisher untersucht werden können, da seine Energie ($\sqrt{s_{NN}} = 5,5$ TeV) die Energie von RHIC ($\sqrt{s_{NN}} = 0,2$ TeV) um ein Vielfaches übersteigt.

Quarkonia

Bei Kollisionen in Teilchenbeschleunigern können schwere Quarks zusammen mit dem entsprechenden Antiquark erzeugt werden. Daraus können sogenannte *Quarkonia* entstehen. Das sind Mesonen, die aus einem Quark und dem entsprechenden Antiquark bestehen: im Fall des Charmonium aus Charm- (c) und Anticharm-Quark ($\bar{c}$), beim Bottonium aus Bottom- (b) und Antibottom-Quark ($\bar{b}$). In dem heißen, dichten QGP wird die Bildung von Quarkonia durch einen Prozess analog zur

Debye'schen Abschirmung bei elektrischer Ladung jedoch unterdrückt [BMS07]. Bei der Hadronisierung bilden c , $\bar{c}$, b und $\bar{b}$ dann mit hoher Wahrscheinlichkeit andere Hadronen (D- und B- Mesonen) mit up- oder down-Quarks. Das führt im Vergleich von Kollisionen, bei denen ein QGP erzeugt wird, zu Kollisionen mit leichten Kernen oder Protonen, bei denen kein QGP entsteht, zu einer Unterdrückung von Quarkonia. Für das J/Ψ , den niedrigsten Zustand des Charmonium, wurde diese auch beim SPS nachgewiesen [NA00].

Auf Grund der höheren Kollisionsenergie beim LHC (ca. das 30-fache bei Pb-Pb-Kollisionen im Vergleich zu RHIC) reicht die Produktion von schwereren Bottom bzw. Antibottom-Quarks erstmals aus, um auch eine Unterdrückung von Υ , dem niedrigsten Zustand des Bottonium, nachzuweisen.

Analog dazu wird erwartet, dass die Produktionsrate von Charm- und Anticharm-Quarks auf ca. 200 Paare pro Kollision ansteigt. Da sich Quarks im QGP frei bewegen können, steigt die Wahrscheinlichkeit, dass Charmonia aus Charm-Quarks von verschiedenen Paaren erzeugt werden, an. Unter Berücksichtigung beider Effekte wird für den LHC in der Summe sogar eine erhöhte Produktion von Charmonia erwartet. [BMS07, YHM05]

Quarkonia zerfallen unter anderem in Elektron-Positron- bzw. Myon-Antimyon-Paare, die bei ALICE nachgewiesen werden können. Diese Elektronen und Positronen werden mit Hilfe des TRD (siehe Kap. 4) identifiziert.

Jet-Quenching

Bei Teilchenkollisionen können harte Streuprozesse mit großem Impulsübertrag zwischen Partonen stattfinden. Diese Partonen fragmentieren dann unmittelbar nach der Streuung im Winkel von 180° jeweils in einen Kegel von Hadronen, die *Jet* genannt werden⁵.

Diese Stöße finden in einer frühen Phase der Kollision statt und die Partonen wechselwirken mit dem QGP. Da die Wechselwirkung in einem QGP sehr viel stärker ist als die in hadronischer Materie, verlieren Partonen, die das QGP durchqueren, sehr viel Energie. Solche Jets unterscheiden sich stark von Jets bei Kollisionen, bei denen kein QGP erzeugt wird. Dieses Phänomen wird *Jet-Quenching* genannt [ALI04].

Auf Grund der großen Anzahl von Teilchen, die bei einer zentralen⁶ Schwerionenkollision erzeugt werden, ist es schwierig Jets zu rekonstruieren. Um sie dennoch zu identifizieren, wird nach führenden Hadronen (*leading particle*) mit hohem transversalen Impuls p_t gesucht. Um die Energie des Jets zu bestimmen, können die Energien der Teilchen, die sich in einem Kegel um das führende Teilchen befinden, integriert werden. Bei RHIC wurde gezeigt, dass Hadronen mit hohem p_t bei zentralen Schwerionenkollisionen im Vergleich zu Protonkollisionen unterdrückt sind, was auf ein QGP hinweist [Ada03].

Wenn ein Jet über ein Hadron mit hohem p_t identifiziert wurde, kann nach dem korrelierten Jet

⁵ Der Winkel von 180° bezieht sich auf das Schwerpunktsystem der beiden Partonen. Im Detektorsystem wird der Winkel von 180° nur in azimuthaler und nicht in z-Richtung beobachtet.

⁶ Die Zentralität einer Kollision ist ein Maß dafür, wieviele Nukleonen der Ionen am Stoß beteiligt sind. Sie ist abhängig vom Abstand der Kernmittelpunkte in der Kollisionsebene.

bei 180° gesucht werden. Auch hier konnte eine Unterdrückung bei Au+Au-Kollisionen am RHIC festgestellt werden. [YHM05, Ada03, BMS07]

Direkte Photonen

Photonen unterliegen nicht der starken Wechselwirkung und werden deshalb nicht wie Jets durch das QGP verändert. Bei einer Kollision gibt es verschiedene Phasen, in denen direkte Photonen entstehen:

- Prompte Photonen, die durch harte Parton–Parton–Streuung unmittelbar bei der Kollision entstehen.
- Thermische Photonen, die während des thermischen Gleichgewichts des QGP durch Kollisionen von Quarks mit anderen Quarks oder Gluonen entstehen.

Diese Photonen enthalten Informationen über die Entwicklung der Kollision. Da sie nicht so stark mit der Materie bzw. dem QGP wechselwirken, dienen sie als Sonde für alle Stadien der Kollision. Thermische Photonen geben z.B. Aufschluss über die Temperatur in der Gleichgewichtsphase.

Zusätzlich zu den direkten Photonen gibt es Zerfallsphotonen, die hauptsächlich aus den Zerfällen von π^0 und η in zwei Photonen stammen. Sie decken den gesamten p_t -Bereich ab und erschweren die Identifizierung direkter Photonen [ALI04]. Direkte Photonen können auch indirekt über den Zerfall in Elektron–Positron–Paare gemessen werden. Diese können mit dem TRD nachgewiesen werden.

Erhöhung der Seltsamkeit

Da die Seltsamkeit (engl. *strangeness*) unter der starken Wechselwirkung eine Erhaltungsgröße ist, können nur $s\bar{s}$ -Paare erzeugt werden.

Der Wirkungsquerschnitt, um in heißer hadronischer Materie aus nicht-seltsamen seltsame Teilchen wie Kaonen und Hyperonen zu erzeugen, ist sehr klein. Dadurch sind große Impulsüberträge notwendig. Der Zerfall $\pi\pi \rightarrow K\bar{K}$ hat zum Beispiel einen Impulsübertrag $Q = 700 \text{ MeV}$ [YHM05]. Außerdem muss eine relativ hohe Mindestenergie erreicht werden, da nur farbneutrale Teilchen und keine einzelnen Quark–Paare produziert werden können. Ein $K\bar{K}$ -Paar, welches das leichteste seltsame Teilchen ist, hat insgesamt eine Ruhemasse von ca. $987 \text{ MeV } c^{-2}$.

Im QGP können $s\bar{s}$ -Paare direkt durch Fusion von zwei Gluonen erzeugt werden. Die Energieschwelle dafür liegt bei ca. 200 MeV [YHM05]. Dadurch ist der Anteil seltsamer Teilchen an allen Teilchen nach der Hadronisierung größer als nach einer Kollision ohne QGP. Dieser Faktor wird *strangeness enhancement factor* genannt.

2.3. Detektion geladener Teilchen

Um geladene Teilchen zu detektieren, müssen sie mit dem Detektor wechselwirken bzw. Energie in ihm deponieren. Deshalb wird in diesem Kapitel zunächst der Energieverlust von Teilchen mit

2. Grundlagen

elektrischer Ladung in Materie, insbesondere die Bethe–Bloch–Formel, diskutiert. Im Anschluss werden zwei Detektortypen, die für diese Arbeit relevant sind, vorgestellt:

- Vieldraht–Proportional–Kammern bzw. Driftkammern
- Szintillatoren in Kombination mit Photoervielfachern

2.3.1. Energieverlust von geladenen Teilchen in Materie

Dringt ein geladenes Teilchen in ein Medium ein, so wechselwirken sie miteinander. Der dominierende Vorgang für den Energieverlust von schweren Projektilen wie Myonen und Protonen sind inelastische Stöße mit den Hüllenelektronen des Materials. Der Energieverlust $\frac{dE}{dx}$ von Teilchen mit der Ladung z in Materie der Ordnungszahl Z und dem Atomgewicht A wird durch die Bethe–Bloch–Formel beschrieben [Ams07]:

$$-\frac{dE}{dx} = \frac{4\pi N_A \alpha^2 z^2}{m_e \beta^2} \left(\frac{Z}{A} \right) \left[\frac{1}{2} \ln \left(\frac{2m_e \beta^2 \gamma^2 T_{max}}{I^2} \right) - \beta^2 - \frac{\delta}{2} \right] \quad (2.3)$$

In der Gleichung ist β die relativistische Geschwindigkeit und γ der Lorentz–Faktor des Teilchens. I ist das effektive Ionisationspotential des Atoms, N_A die Avogadro–Konstante, m_e die Elektronenmasse, α die Feinstrukturkonstante und δ eine kleine Korrektur, die bei hohen Elektronendichten

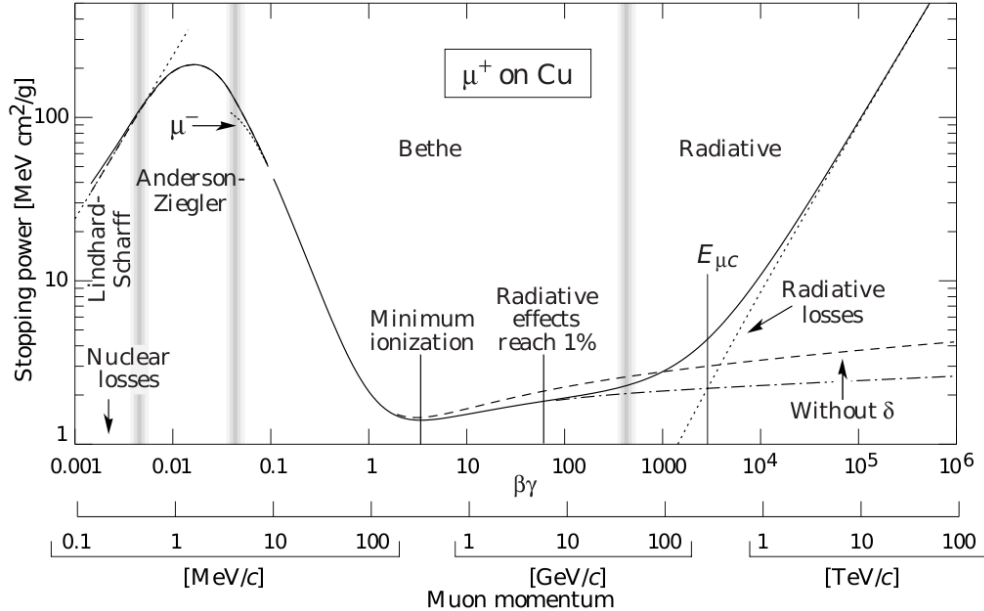


Abb. 2.4.: Der Energieverlust eines Myons in Kupfer pro Längeneinheit in Abhängigkeit des Impulses. Die durchgezogene Kurve zeigt den gesamten Energieverlust. Für mittlere Impulse ist die Bethe–Bloch–Formel eine gute Näherung. Für $E > E_{\mu c}$ dominieren radiative Effekte. Myonen die im Minimum bei $\beta\gamma = 5$ liegen, werden minimal ionisierende Teilchen genannt, da ihr Energieverlust minimal ist. Bild aus [PDG10].

und hohen Energien nötig ist. T_{max} ist der maximale Energieübertrag an ein Elektron. Die Bethe–Bloch–Formel hat für schwere Projektile im Bereich $0,1 \lesssim \beta\gamma \lesssim 1000$ eine Genauigkeit von wenigen Prozent [PDG10].

In Abb. 2.4 ist der spezifische Energieverlust eines Myons in Kupfer in Abhängigkeit des Impulses aufgetragen. Der Bereich, in dem die Bethe–Bloch–Formel als Näherung gültig ist, ist durch die senkrechten Striche markiert. Der Energieverlust findet in diesem Bereich hauptsächlich durch Stöße mit den Hüllenelektronen statt. Teilchen, die im Minimum der Kurve liegen, werden minimal ionisierende Teilchen genannt. Bei niedrigeren Energien sind die Geschwindigkeiten des Projektils und der Hüllenelektronen vergleichbar und verschiedene Korrekturen müssen berücksichtigt werden. Bei hohen Energien dominieren radiative Effekte. Für Elektronen und Positronen muss Bremsstrahlung auf Grund ihres Gewichts schon bei geringeren Energien berücksichtigt werden.

2.3.2. Vieldraht–Proportional–Kammern

Mit Vieldraht–Proportional–Kammern (MWPC kurz für *Multi Wire Proportional Chamber*) können die Koordinaten, an denen ein geladenes Teilchen die Kammer durchquert, bestimmt werden. In Abb. 2.5 ist das Funktionsprinzip einer Vieldraht–Proportional–Kammer dargestellt. Die Kammer besteht aus einem geschlossenen Volumen, in dem sich ein leicht ionisierbares Gas befindet. An die Kathodendrähte wird eine negative Spannung im Bereich von einigen Kilovolt angelegt, während die Anodendrähte auf Erdpotenzial gehalten werden.

Einfallende Teilchen ionisieren das Gas und die Elektronen, die dabei herausgelöst werden, wandern in dem elektrischen Feld zu den Anodendrähten. In der Nähe der Drähte werden Elektronen durch das inhomogene Feld ($\propto 1/r$) stark beschleunigt und ionisieren so lawinenartig weitere Atome. Die Verstärkung liegt typischerweise in der Größenordnung 10^4 , wobei das erzeugte Signal proportional zu der Energie ist, die das Teilchen in der Kammer deponiert hat. Die Ortsauflösung liegt bei einem Anodendrahtabstand $d = 2 \text{ mm}$ typischerweise bei $\sigma = 600 \text{ } \mu\text{m}$. [Ams07]

Wird eine in Streifen segmentierte Folie als Kathode benutzt und die induzierte Ladung der Ionen ausgelesen, kann die Koordinate senkrecht zur Anodenkoordinate bestimmt werden.

In einer Driftkammer wird zusätzlich noch die Driftzeit der Elektronen gemessen, die bei einem homogenen Driftfeld proportional zu der Position ist, an der es aus dem Atom ausgelöst wurde. So können nicht nur die Koordinaten, an der das Teilchen die Kammer durchquert hat, sondern auch die dreidimensionale Trajektorie bestimmt werden. Befindet sich die Kammer in einem Magnetfeld, werden die Teilchen auf gekrümmte Bahnen gezwungen und der Impuls der Teilchen kann anhand

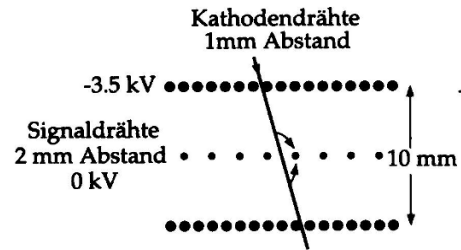


Abb. 2.5.: Das Funktionsprinzip einer Vieldrahtkammer. Geladene Teilchen ionisieren das Gas in der Kammer. Die herausgelösten Elektronen werden verstärkt und ausgelesen. Bild aus [Ams07]

der Krümmung bestimmt werden.

Da der TRD auch eine spezielle Driftkammer ist, wird das Prinzip in Kap. 4.4 wieder aufgegriffen und speziell für den TRD erläutert.

2.3.3. Szintillatoren und Photovervielfacher

Szintillatoren sind aus Materialien gefertigt, die einen Teil der Energie eines geladenen Teilchens oder Photons, von dem sie durchquert werden, absorbieren und in Form eines Photons wieder abgeben. Es wird zwischen anorganischen und organischen Szintillatoren unterschieden.

Anorganische Szintillatoren

Anorganische Szintillatoren bestehen aus Einkristallen wie NaI oder CsI. Einfallende Teilchen erzeugen durch Ionisation Elektron–Loch–Paare, bei deren Vernichtung Licht emittiert wird. Anorganische Szintillatoren eignen sich aufgrund ihrer hohen Ordnungszahl gut zum Nachweis von Photonen. Allerdings haben sie eine relativ lange Abklingzeit (20 ns bei CsI) und sind auf Grund der Einkristalle relativ teuer. [Ams07]

Organische Szintillatoren

Organische Szintillatoren bestehen üblicherweise aus Plastik und werden deshalb auch Plastik–Szintillatoren genannt. Dem Plastik ist ein primärer Fluoreszenzstoff beigemischt, in dem durch einfallende Teilchen molekulare Zustände angeregt werden, die beim Zerfall Photonen im UV–Bereich emittieren. Darüber hinaus wird in der Regel ein Wellenlängenschieber in geringer Konzentration beigemischt, der die erzeugten Photonen absorbiert und dann Photonen im sichtbaren Bereich emittiert. Indem die Wellenlänge verschoben wird, wird die Reabsorption der Photonen im Szintillator unterdrückt.

Organische Szintillatoren eignen sich wegen ihrer kleinen Ordnungszahl schlecht zur Detektion von Photonen, haben dafür aber schnelle Abklingzeiten von ca. 2 ns und können großflächig und günstig hergestellt werden. [Ams07]

Für das Triggersystem in Münster, das Gegenstand dieser Arbeit ist, werden auch Plastik–Szintillatoren verwendet, die in Kap. 6.2 vorgestellt werden.

Photovervielfacher

Die Photonen, die mit Szintillatoren erzeugt werden, werden in der Regel mit einem Photovervielfacher (PMT kurz für *PhotoMultiplier Tube*) detektiert, verstärkt und als elektrisches Signal ausgegeben.

In Abb. 2.6 ist die Funktionsweise einer PMT schematisch dargestellt. Photonen treffen auf die Photokathode, wo sie durch den Photoeffekt Elektronen auslösen. Je nach Modell liegt an der

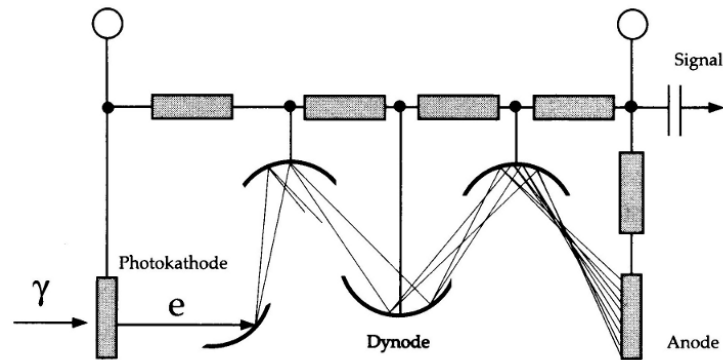


Abb. 2.6.: Das Funktionsprinzip eines Photovervielfachers. In der Photokathode werden durch die einfallenden Photonen Elektronen emittiert, die dann über Dynoden vervielfacht und an der Anode ausgelesen werden.

Kathode eine negative oder an der Anode eine positive Spannung im kV-Bereich an, die über Spannungsteiler auf die Dynoden geteilt wird. Die Elektronen werden durch das Spannungspotential von Dynode zu Dynode beschleunigt. Treffen sie auf eine Dynode, werden durch den Aufprall Sekundärelektronen emittiert, die dann zur nächsten Dynode beschleunigt werden. Das führt zu einer lawinenartigen Verstärkung der Primärelektronen. An der Anode wird die erzeugte Ladung ausgelesen und in der Regel als Spannungspuls weitergegeben. Die Anzahl der Dynoden, die Hochspannung und somit auch die Verstärkung hängen vom jeweiligen Modell ab. Typischerweise liegt sie in der Größenordnung 10^8 . [Ams07]

Die Photovervielfacher, die für das Triggersystem in Münster verwendet werden, werden in Kap. 6.2 vorgestellt.

In diesem Kapitel wurden die Grundlagen für diese Arbeit diskutiert. Zunächst wurde eine kurze historische Einführung in die Teilchenphysik gegeben und das Standardmodell der Elementarteilchen vorgestellt. Daraufhin wurde über die starke Wechselwirkung ein Bezug zum Quark-Gluon Plasma hergestellt und diskutiert, dass mit Schwerionenkollisionen in Teilchenbeschleunigern ein QGP erzeugt und durch charakteristische Signaturen nachgewiesen und untersucht werden kann.

Zum Schluss wurden die Grundlagen zur Detektion von geladenen Teilchen diskutiert und das Prinzip von Vieldraht- bzw. Driftkammern erläutert, da der TRD diese Technik benutzt. Außerdem wurden Szintillatoren und wie sie mit Photovervielfachern ausgelesen werden, vorgestellt. Diese werden benutzt, um kosmische Teilchen nachzuweisen, mit denen die Supermodule in Münster kalibriert werden können.

3. Der LHC und ALICE

Im ersten Unterkapitel 3.1 wird der Large Hadron Collider mit einigen seiner Kenngrößen und sein Funktionsprinzip erklärt. Am Beschleuniger befinden sich vier große Experimente. Eins davon ist A Large Ion Collider Experiment, das im Unterkapitel 3.2 vorgestellt wird. Dazu werden die einzelnen Subdetektoren kurz erklärt und ein besonderes Augenmerk auf das komplexe Trigger-System von ALICE gelegt.

3.1. Large Hadron Collider

Der LHC ist mit einem Umfang von ca. 27 km der größte Teilchenbeschleuniger, der je gebaut wurde. Er wurde 1996 vom CERN Council bewilligt und im Herbst 2008 zum ersten Mal in Betrieb genommen. Durch ein massives Kühlungsproblem musste er aber schon nach wenigen Tagen wieder abgeschaltet werden. Um solchen Problemen in Zukunft vorzubeugen, wurde das gesamte Kühlungssystem gewartet. Ende 2009 wurde der LHC dann wieder in Betrieb genommen. Bis November 2010 standen nur Proton-Proton-Kollisionen auf der Agenda. Dazu wurde die Energie zunächst sukzessive auf $\sqrt{s_{NN}} = 7$ TeV, was der Hälfte der maximalen Energie von $\sqrt{s_{NN}} = 14$ TeV für Proton-Kollisionen entspricht, erhöht.

Seit November 2010 werden erste Bleikollisionen durchgeführt. Die maximale Schwerpunktennergie bei Bleikollisionen ist 1146 TeV, was einer Energie von $\sqrt{s_{NN}} = 5,5$ TeV pro Nukleonpaar entspricht. Damit überragt der LHC den bisher stärksten Schwerionen-Beschleuniger RHIC ($\sqrt{s_{NN}} = 0,2$ TeV) um das 30-fache. Durch die hohe Luminosität¹ von $1 \times 10^{27} \frac{1}{\text{cm}^2 \text{ s}}$ [Brü04] wird außerdem eine hohe Statistik erreicht.

Am LHC gibt es vier große Experimente: ALICE, *A Toroidal LHC ApparatuS* (ATLAS), *Compact Muon Solenoid* (CMS) und *Large Hadron Collider beauty* (LHCb). ATLAS, CMS und LHCb sind für die Untersuchung von Proton-Proton-Kollisionen ausgelegt und sollen unter Anderem supersymmetrische Teilchen und das theoretisch vorhergesagte Higgs-Boson, das Teilchen Masse verleiht, nachweisen. ALICE ist das einzige Experiment, das für die Kollisionen von Bleikernen optimiert ist, um ein QGP zu untersuchen. Es wird in Kapitel 3.2 noch detailliert beschrieben.

Der LHC wurde in den ehemaligen Tunnel des *Large Electron-Positron Collider* (LEP) gebaut. Dieser befindet sich im Mittel ca. 100 m unter der Erde in der Nähe von Genf. Darin sind zwei

¹ Die Luminosität ist ein Maß für die Anzahl von Teilchen (in diesem Fall Bleiionen) pro Zeit- und Flächeneinheit. Zusammen mit dem Wirkungsquerschnitt lässt sich daraus die erwartete Anzahl produzierter Teilchen berechnen.

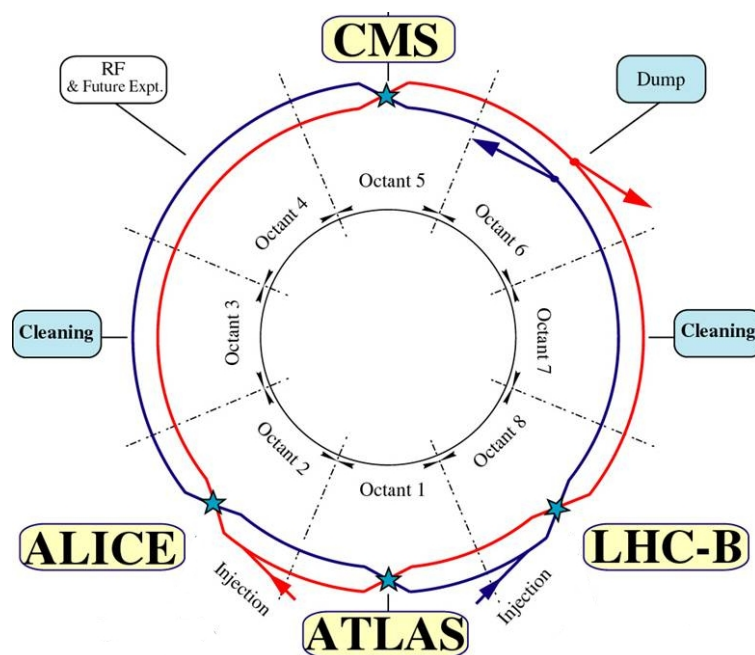


Abb. 3.1.: Die Einteilung des LHC in acht Segmente. Es sind die acht Interaktionspunkte vorhanden. Davon sind vier Kollisionspunkte, an denen die großen Experimente ALICE, ATLAS, CMS und LHCb gelegen sind. Bild aus [CER10].

27 km lange Strahlrohre installiert, in denen Hadronen einmal gegen und einmal im Uhrzeigersinn beschleunigt werden. In den Strahlrohren herrscht ein Ultrahochvakuum von 10^{-13} atm [CER08], um sicherzustellen, dass die zirkulierenden Teilchen nicht mit Restmolekülen stoßen. In Abb. 3.1 ist der LHC mit seinen Experimenten schematisch dargestellt. Er stellt keinen perfekten Kreis dar, sondern besteht aus acht Bögen und acht geraden Strecken, die als Kollisionspunkte benutzt werden oder um den Strahl zu entsorgen oder zu manipulieren.

An den Kollisionspunkten kreuzen sich die beiden Teilchenstrahlen und Kollisionen können mit den Detektoren beobachtet werden.

Die Protonen bzw. Bleikerne werden durch supraleitende Hohlraumresonatoren beschleunigt, bis sie ihre maximale Energie erreicht haben. Durch Resonanzüberhöhung werden mit Radiofrequenzen in den Kavitäten Felder erzeugt, die viel größer als die der eigentlichen elektromagnetischen Wellen sind. Damit die Teilchen ein Potential, das nur beschleunigend ist, erfahren, muss sich das Vorzeichen des Feldes beim Durchqueren umdrehen. Daraus resultiert auch unmittelbar, dass im LHC keine kontinuierlichen Teilchenstrahlen, sondern Teilchenpakete, sogenannte *Bunches*, zirkulieren. Je Strahlrohr sind acht dieser Hohlraumresonatoren, die ein Feld von bis zu 5 MV/m mit einer Frequenz von 400 MHz [CER08] erzeugen, installiert. Sie müssen auf 4,5 K gekühlt werden. Neben dem Beschleunigen der Teilchen besteht ihre Hauptaufgabe darin, die Bunches möglichst kompakt zu halten, um somit möglichst hohe Kollisionsraten bei den Experimenten zu erzielen. Bei Blei-Kollisionen sind 592 Bunches mit jeweils 7×10^7 Bleionen im LHC enthalten. Damit soll eine

Luminosität von $1 \times 10^{27} \frac{1}{\text{cm}^2 \text{ s}}$ [Brü04] erreicht werden.

Damit Hadronen in den Strahlrohren zirkulieren, muss ihre Richtung durch Magnete manipuliert werden. Dies geschieht durch 9593 supraleitende Dipol- und Multipol-Magnete, die bei 1,9 K betrieben werden. Indem supraleitende Magnete verwendet werden, können durch die Dipolmagnete 11700 A fließen und es werden Feldstärken von 8,3 T erzeugt. [CER08]

Weil die Magnete und die Hohlraumresonatoren mit flüssigem Helium auf 1,9 K bzw. 4,5 K gekühlt werden müssen, musste ein beispielloses Kühlsystem entwickelt werden, das eine Kühlkapazität von insgesamt 170 kW hat. Es dauert mehrere Wochen, um die Magneten auf 1,9 K zu kühlen oder sie wieder aufzuwärmen. Das ist zum Beispiel notwendig, wenn der LHC abgeschaltet wird, um Detektoren oder den Beschleuniger zu warten oder zu erweitern. [CER08]

Um Bleikerne bzw. Protonen in den LHC einzuspeisen und die Energien von $\sqrt{s_{NN}} = 5,5 \text{ TeV}$ (14 TeV) bei Bleikernen (Protonen) zu erreichen, werden die Teilchen zunächst in diversen anderen Beschleunigern vorbeschleunigt. In Abb. 3.2 ist dies schematisch dargestellt und im Detail erläutert.

3.2. ALICE

A Large Ion Collider Experiment ist für Bleikollisionen optimiert, bei denen ein QGP nachgewiesen und untersucht werden soll. Da das QGP jedoch nur für ca. 10^{-22} s [BMS07] existiert, kann es nur über die Teilchen, die bei der Kollision produziert werden, nachgewiesen werden. Dazu wurden in Kapitel 2.2 bereits einige typische Signaturen aufgelistet, anhand derer ein QGP identifiziert werden kann. Da mit dem LHC Schwerpunktennergien von 1148 TeV bei Bleikollisionen erreicht werden, wird angenommen, dass das QGP mit ALICE genauer untersucht werden kann als mit allen bisherigen Experimenten.

Wie viele moderne Detektoren besteht ALICE aus einzelnen Subdetektoren, die zusammen ca. 3.000 [ALI10a] geladene Teilchen, die bei einer Kollision² entstehen, identifizieren und ihre Trajektorie vermessen sollen. In Abb. 3.3 ist der Detektor schematisch dargestellt. Einige charakteristische Eigenschaften und Aufgaben der einzelnen Subdetektoren werden im Folgenden beschrieben.

3.2.1. Die Subdetektoren von ALICE

Der ALICE-Detektor besteht aus einem Zylinder, der symmetrisch um den Kollisionspunkt längs der Flugachse der Ionen platziert ist, und dem *Myon-Spektrometer* in Vorwärts-Richtung. Der Zylinder ist von einem 12 m langen Magneten umgeben, der im Durchmesser 10–12 m misst und ein homogenes Magnetfeld von 0.5 T erzeugt. Dadurch werden Teilchen, die eine elektrische Ladung tragen, je nach Vorzeichen im bzw. gegen den Uhrzeigersinn auf Helixbahnen gelenkt. Dies erlaubt es mit verschiedenen Detektoren, die an der Spurrekonstruktion beteiligt sind, den Impuls der Teilchen zu bestimmen.

² Pb–Pb-Kollisionen mit halber Energie, also $\sqrt{s_{NN}} = 2,76 \text{ TeV}$.

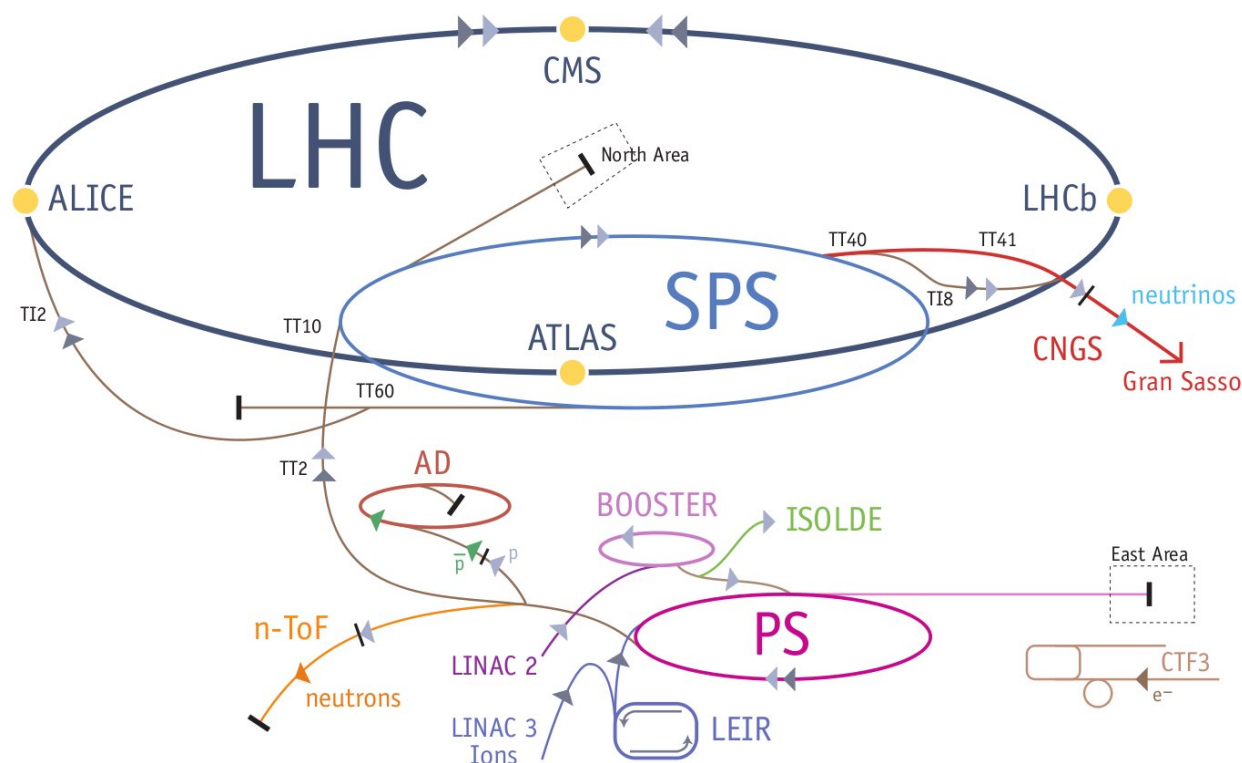


Abb. 3.2.: Blei-Ionen werden erzeugt, indem zunächst eine sehr reine Bleiprobe auf ca. 550 °C erhitzt wird. Die Blei-Atome im Dampf werden dann durch Beschuss mit schnellen Elektronen bis zu Pb^{29+} ionisiert. Diese werden dann auf $4,2 \text{ MeV u}^{-1}$ (Energie pro Nukleon) beschleunigt und mit einer Kohlenstoff-Folie weiter zu Pb^{54+} ionisiert, um in den Low Energy Ion Ring (LEIR) injiziert und auf 72 MeV u^{-1} beschleunigt zu werden. Im Proton Synchrotron (PS) erreichen sie dann $5,9 \text{ GeV u}^{-1}$ und werden anschließend durch eine zweite Stripper-Folie vollständig zu Pb^{82+} ionisiert. Bevor sie in den LHC injiziert werden, werden sie im Super Proton Synchrotron (SPS) noch auf 177 GeV u^{-1} beschleunigt. Im LHC können dann $2,76 \text{ TeV u}^{-1}$ erreicht werden. Im November 2010 wurden erste Bleikollisionen mit halber Energie ($1,38 \text{ TeV u}^{-1}$) durchgeführt. Bei Protonkollisionen wurde bisher eine Energie von $3,5 \text{ TeV}$ im LHC erreicht (Stand: Februar 2011). Diese werden aus Wasserstoff gewonnen und ebenfalls über mehrere Stationen beschleunigt. [CER08].

Direkt um das Strahlrohr ist das *Inner Tracking System* (ITS) angebracht. Es besteht aus sechs Lagen von hochauflösenden Halbleiter-Detektoren. Die Aufgaben des ITS umfassen die Lokalisierung des primären Vertex und der sekundären Vertices von schnell zerfallenden Teilchen, die seltsame oder schwerere Quarks beinhalten. Außerdem verbessert es die Impuls-Messung der Teilchen.

An das ITS schließt sich die *Time Projection Chamber* (TPC) an. Ihr Radius beträgt zwischen 88 cm und 250 cm und hat ein Volumen von 88 m^3 . Die TPC wird mit Vieldraht-Proportional-Kammern ausgelesen und ist der Haupt-Detektor für die Spurrekonstruktion. Dadurch produziert sie sehr viele Daten (ca. 60 MB/Ereignis) und wird nur ausgelesen, wenn schnellere Detektoren

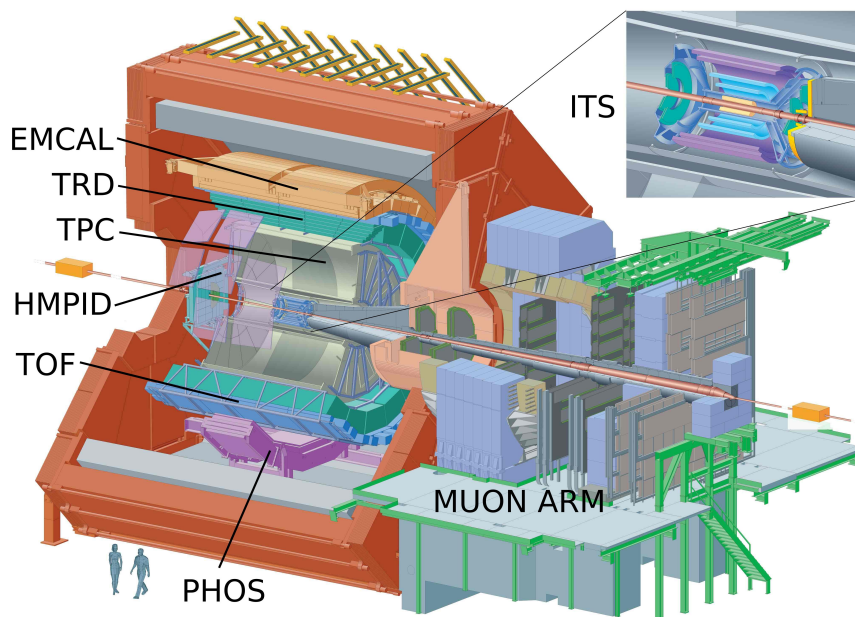


Abb. 3.3.: Der ALICE Detektor mit seinen Subdetektoren. Er besteht aus einem zentralen Zylinder und dem Myon-Arm in Vorwärts-Richtung. In rot ist der große L3-Magnet zu erkennen, mit dem elektrisch geladene Teilchen abgelenkt werden. Bild aus [Kle08]

das Ereignis als relevant einstufen³. Um die TPC ist der *Transition Radiation Detector* (TRD) installiert, der Inhalt dieser Arbeit ist und deshalb detailliert in Kap. 4 besprochen wird. Nach dem TRD ist der *Time Of Flight-Detektor* (TOF) verbaut. Er ist in der Lage die Flugzeiten der Teilchen vom Vertex mit einer Genauigkeit besser als 100 ps zu bestimmen. Sowohl ITS, TPC, TRD als auch TOF decken eine Pseudo-Rapidity⁴ im Bereich $|\eta| < 0,9$ ab [ALI04].

Zusätzlich sind außerhalb des Zylinders noch einige Detektoren, die nur kleinere Winkelbereiche abdecken, angebracht. Der *High Momentum Particle Identification Detector* (HMPID) dient der Identifikation von Teilchen mit hohem transversalen Impuls p_t . Er besteht aus sieben *Ring Imaging Cherenkov* (RICH) – Detektormodulen. Er deckt eine Pseudo-Rapidity von $|\eta| < 0,6$ und einen azimuthalen Winkel von ca. 58° ab. Um Photonen, insbesondere direkte Photonen, mit hoher Winkel- und Energie-Auflösung zu identifizieren, wurde das *PHOTon Spectrometer* (PHOS) integriert. Es besteht aus 5 Modulen, die sich jeweils aus einem fein segmentierten elektromagnetischen Kalorimeter in Kombination mit einem Detektor (CPV kurz für *Charged Particle Veto*), mit dem geladene Teilchen von der Messung ausgeschlossen werden, zusammensetzen. Es kann Beiträge zu verschiedenen Triggerentscheidungen (Level 0 und Level 1) liefern und deckt $|\eta| < 0,12$ und einen Azimutalwinkel von 100° ab [ALI04]. Photonen können außerdem über das *ElectroMagnetic Calorimeter* (EMCal) detektiert werden, das insbesondere zur Jet-Detektion genutzt wird und auch die Möglichkeit eines

³ Die TPC wird nur nach einem Level 1-Trigger ausgelesen. Das Triggersystem von ALICE wird in Kap. 3.3 noch ausführlich besprochen.

⁴ Die Pseudo-Rapidity $\eta = -\ln(\tan(\theta/2))$ ist ein Maß für den Winkel relativ zur Strahlachse

Jet-Triggers bietet [ALI08].

Es gibt noch weitere kleinere Detektoren, die für spezielle Aufgaben zu ALICE hinzugefügt wurden. Mit *Photon Multiplicity Detector* (PMD), *Forward Multiplicity Detector* (FMD) und *Zero Degree Calorimeter* (ZDC) lassen sich Aussagen über globale Eigenschaften einer Kollision, wie Multiplizität und Zentralität machen. Mit *T0* und *V0* wird ein schneller Pretrigger für den TRD generiert, um die Elektronik auf ein Ereignis vorzubereiten. T0 generiert außerdem einen Level 0-Trigger-Beitrag und das Startsignal für TOF. V0 gibt Informationen über den Vertex und die Zentralität einer Kollision und liefert dementsprechend einen Trigger. T0 und V0 werden noch im Rahmen des Pretrigger-Systems detaillierter in Kap. 4.5 diskutiert. Über dem L3 Magneten befindet sich *A COsmic Ray Detector for ALICE* (ACORDE), ein Verbund von Szintillatoren, der als Trigger für kosmische Teilchen benutzt werden kann, um ALICE zu testen, wenn der LHC nicht in Betrieb ist.

In Vorwärts-Richtung ist der Myon-Arm, dessen Hauptaufgabe darin besteht, Myonen bzw. Myonen-Paare nachzuweisen, die aus dem Zerfall von Quarkonia hervorgehen. Noch innerhalb des L3-Magneten befindet sich ein Absorber, der Hadronen und Photonen abschirmt. Das Myon-Spektrometer besteht aus einem eigenen Dipol-Magneten, Spurrekonstruktions- und Triggerdetektoren, die eine Pseudo-Rapidity von $-4,5 \leq \eta \leq -2,5$ abdecken. Vor den Triggerdetektoren ist noch ein zusätzlicher Myon-Filter eingebaut, um diese zu schützen. Durch die Filter können nur Myonen mit einem Impuls größer als 4 GeV/c nachgewiesen werden.

In Abb. 3.4 ist eine Rekonstruktion einer der ersten Blei-Blei-Kollisionen vom 8.11.2010 gezeigt, die mit ALICE aufgezeichnet wurden.

3.3. Das ALICE-Triggersystem

In diesem Kapitel soll das globale ALICE-Triggersystem vorgestellt werden. Bei der angestrebten Luminosität wird eine Kollisionsrate von ca. 8 kHz erwartet. Die TPC, der langsamste Detektor mit dem größten Datenvolumen, ist bei zentralen Pb-Pb-Kollisionen aufgrund seiner Auslesezeit auf ca. 200 Hz begrenzt [ALI04]. Dadurch muss aus den zur Verfügung stehenden Ereignissen ausgewählt werden.

Grundlage für die zu untersuchende Physik ist eine hohe Statistik. Ziel ist es, möglichst viele seltene Ereignisse aufzuzeichnen, während häufige, die unter Umständen schon weitestgehend untersucht wurden, zu vernachlässigen sind. Mit dieser Vorgabe wurde ein Triggersystem entwickelt, das nur unter bestimmten Bedingungen die Detektoren bzw. ihre Auslese aktiviert oder ggf. abbricht und bereits ausgelesene Daten verwirft. Mit diesem ist es möglich das generierte Datenvolumen erheblich zu reduzieren und die Rate an die langsamen Detektoren anzupassen, ohne Statistik bei seltenen Ereignissen zu verlieren.

Um den verschiedenen Auslesegeschwindigkeiten der Detektoren gerecht zu werden, wurde ein mehrstufiges Trigger-System entwickelt. Es gibt fünf verschiedene Triggerlevel, für die jeweils die

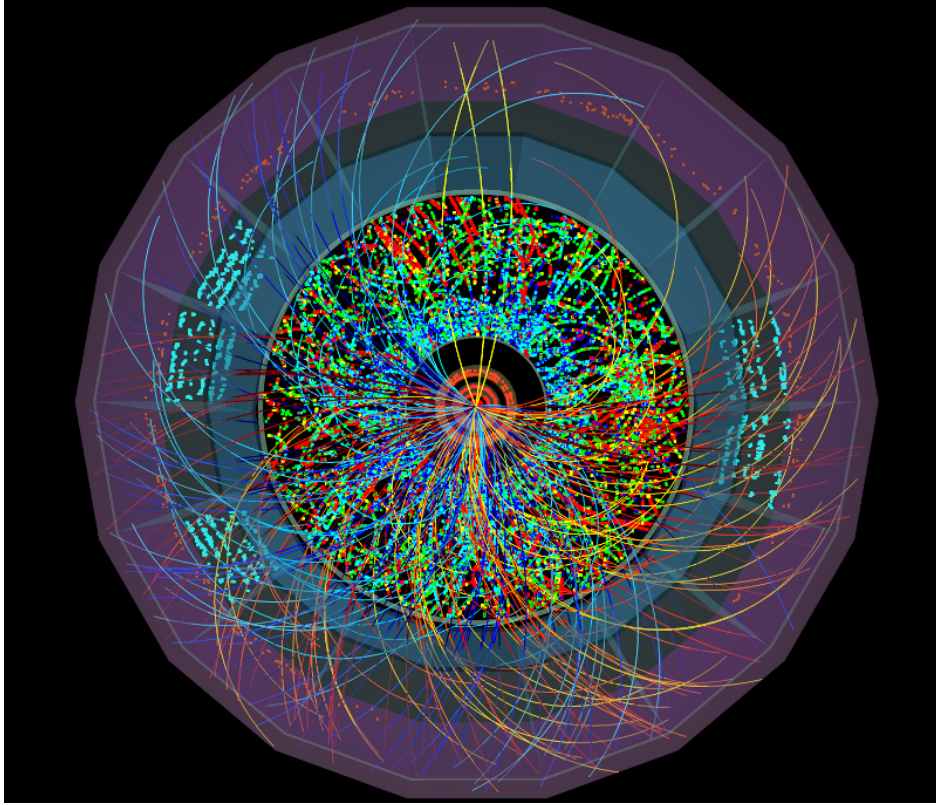


Abb. 3.4.: Eines der ersten Pb–Pb–Ereignisse in ALICE vom 8.11.2010. Durch die Farben ist die Energie der Teilchen angedeutet. Die schattierten Elemente sind die einzelnen Detektoren. Bild aus [ALI10]

zu dem Zeitpunkt zur Verfügung stehenden Informationen der einzelnen Detektoren ausgewertet werden: *Pretrigger* (PT), *Level 0* (L0), *Level 1* (L1), *Level 2* (L2) und der *High Level Trigger* (HLT).

In Abb. 3.5 ist die Triggerhierarchie dargestellt. Der PT ist ein schneller Hardware-Trigger, der ausschließlich dazu dient, die Elektronik des TRD aufzuwecken. Durch die Implementation des Pretriggers kann die Elektronik die meiste Zeit in einem Standby-Modus gehalten werden. L0, L1 und L2 sind Hardware-Trigger und werden von dem *Central Trigger Processor* (CTP) generiert, indem insgesamt 50 Triggerbeiträge der Subdetektoren ausgewertet werden. Der HLT ist ein Software-Trigger, der eine Online-Rekonstruktion der Kollisionen vornimmt. Nach dem HLT werden die Ereignisse dem *Data Acquisition-System* (DAQ) zugeführt und dauerhaft gespeichert. Die Trigger werden über das *Trigger, Timing and Control* (TTC) – System an die einzelnen Subdetektoren verteilt. Im Folgenden sollen die einzelnen Triggerlevel, der CTP und das Distributionsnetzwerk vorgestellt werden.

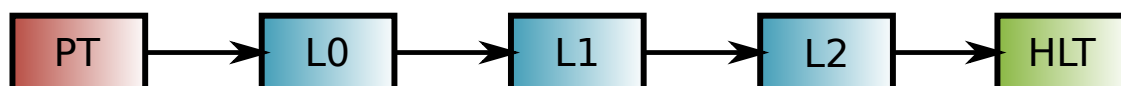


Abb. 3.5.: Der Ablauf einer vollständigen Triggersequenz. Pretrigger, Level 0, Level 1 und Level 2 sind Hardware-Trigger. Der PT ist ein schneller Trigger, der lediglich dazu dient die Elektronik des TRD aufzuwecken. L0 – L2 werden von der CTP durch insgesamt 50 Triggerbeiträge der Subdetektoren generiert. Der HLT ist ein Software-Trigger, bei dem eine Kollision mit Hilfe einer Computer-Farm online, d.h. während Kollisionen stattfinden, rekonstruiert und analysiert werden.

3.3.1. Pretrigger

Der Pretrigger ist ein schneller Hardwaretrigger. Auf den TRD Modulen ist massiv parallele Elektronik implementiert (siehe Kap. 4.6), um schon $6,1 \mu\text{s}$ nach der Kollision einen L1-Triggerbeitrag zu liefern. Die Elektronik wird nur aktiviert, wenn sie auch wirklich benötigt wird.

Damit sie rechtzeitig für ein Ereignis bereit ist, wurde das Pretrigger-System integriert, das mit einer geringen Latenz von ca. 716 ns [PTS10] nach der Kollision einen Trigger bereitstellt.

Eine Kopie des PT wird als L0-Beitrag an den CTP weitergegeben. Da das Pretrigger-System für den TRD implementiert wurde, wird es in Kap. 4.5 detailliert diskutiert.

3.3.2. Central Trigger Processor: Level 0, Level 1, Level 2

Der CTP ist zuständig für die Trigger Level 0, Level 1 und Level 2. Um sie zu generieren, verarbeitet der CTP insgesamt 50 Triggerbeiträge. In Anhang A sind alle Triggerbeiträge, die zur Verfügung stehen, aufgelistet. Innerhalb des CTP werden 50 Klassen, in denen Triggerbeiträge für die verschiedenen Level zusammengefasst werden, definiert.

Die Subdetektoren von ALICE werden in sechs dynamische Cluster eingeteilt, die unabhängig voneinander getriggert werden können. Dadurch können Gruppen von Subdetektoren mit verschiedenen Raten getriggert bzw. ausgelesen werden.

Damit nur Trigger gesendet werden, wenn alle beteiligten Systeme bereit sind, geben die Subdetektoren und die DAQ jeweils ein *busy*-Signal zurück, das das jeweilige System als beschäftigt markiert.

Um zu vermeiden, dass sich Ereignisse bei den Detektoren anstauen oder überschneiden, sind für jeden Level vier Vergangenheits-Zukunfts-Absicherungen (PFP kurz für *past-future-protections*) implementiert. Sie verwerfen einen Trigger, wenn in einem frei definierbaren Zeitfenster um den Trigger ein anderer Trigger generiert werden soll.

Cluster

Die Subdetektoren haben verschiedene Sensitivitätszeiten bzw. Auslesezeiten. Um die Statistik zu erhöhen, ist es in einigen Situationen sinnvoll, verschiedene Subdetektoren mit unterschiedlichen Raten auszulesen bzw. zu triggern.

Die Logik, die benötigt würde, um jeden Subdetektor einzeln zu triggern, ist jedoch zu umfangreich. Deshalb wurden sechs dynamische Cluster eingerichtet, in denen jeweils eine Gruppe von Subdetektoren zusammengefasst wird. Jeder Cluster kann unabhängig von den anderen Clustern getriggert werden.

Ein Cluster hat eine oder mehrere Klassen als Eingang, die die Triggerbeiträge verarbeiten. Ein Cluster wird getriggert, wenn eine der Klassen, die mit ihm verknüpft ist, aktiv ist.

Klassen

Die CTP verarbeitet insgesamt 50 Triggerbeiträge (siehe Anhang A). Nicht jede mögliche Kombination aus Triggerbeiträgen ist dabei physikalisch sinnvoll. Da darüber hinaus auch hier die Logik zu umfangreich wäre, um alle Kombinationen abzudecken, werden 50 Klassen definiert.

Eine Klasse fasst mehrere Triggerbeiträge für jeden Level zusammen. Das bedeutet, dass in einer Klasse definiert ist, welche Triggerbeiträge jeweils aktiv sein müssen, damit ein L0, L1 bzw. L2 ausgegeben wird. Die definierten Beiträge werden separat für jeden Level mit einem logischen UND verschaltet. In Anhang A ist eine Liste der verfügbaren Klassen aufgelistet.

Jede Klasse ist mit genau einem Cluster verknüpft und gibt die Trigger-Entscheidung an ihn weiter. Sind z.B. alle Beiträge für den L0 aktiv, wird ein L0 für den entsprechenden Cluster ausgegeben.

Bevor eine Klasse einen Trigger an einen Cluster weitergibt, werden noch verschiedene Veto-Signale geprüft, die den Trigger gegebenenfalls verhindern. Da die Vetos für die verschiedenen Level unterschiedlich sind, werden sie an entsprechender Stelle erläutert.

Level 0

Insgesamt hat der CTP 24 L0-Eingänge (siehe Anhang A), die jeder Klasse zur Verfügung stehen. Welche der Beiträge für die Klasse genutzt werden, kann jedoch frei konfiguriert werden. Der L0-Trigger ist ein 25 ns langer Puls, der spätestens $1,2 \mu\text{s}$ nach der Kollision an den Subdetektoren erwartet wird. Trifft er dort in diesem Zeitfenster ein, wird er als *L0-Accept* (L0A), also als gültiger L0 gewertet. Bleibt der Puls aus, wird die Auslese abgebrochen bzw. erst gar nicht gestartet (L0R kurz für *L0-Reject*).

Für den L0 stehen verschiedene Vetos zur Verfügung, die verhindern, dass ein L0 gesendet wird. Verpflichtend für jede Klasse sind DAQ-busy, CTP-busy, CTP-Totzeit und Cluster-busy. Für das Cluster-busy werden alle busy-Signale von den Detektoren des Clusters, mit dem die Klasse verknüpft ist, mit einem logischen ODER verknüpft. Vier past-future-protections, die verhindern, dass sich Ereignisse anstauen oder überlappen, können frei eingestellt und als Veto benutzt werden.

Es dürfen nur neue Triggersequenzen gestartet werden, wenn keines der Vetos aktiv ist. Zwischen L0 und L1 sind neue Sequenzen verboten, da dann mindestens das Cluster-busy aktiv ist. In der langen Zeit zwischen L1 und L2 sind sie jedoch erlaubt.

Level 1

Nach $6,5 \mu\text{s}$ wird die L1-Entscheidung an den Subdetektoren erwartet, die auf Grund von 20 Beiträgen getroffen wird. Der L1-Trigger besteht aus einem 50 ns langem Puls und einer L1-Triggernachricht (L1M). Analog zum L0 wird der L1-Puls in einem festen Zeitfenster erwartet und wird dementsprechend als L1A bzw. L1R interpretiert. Kurz nach dem L1A wird die L1M gesendet, die Details zu dem Trigger beinhaltet (siehe auch Kap. 3.3.4).

Als Veto sind beim L1 nur die vier past-future-protections vorgesehen, die auch hier frei konfiguriert und gewählt werden können.

Level 2

Nach $88 \mu\text{s}$ ist die Driftzeit der TPC beendet und der CTP trifft mit sechs L2-Beiträgen eine Level 2-Entscheidung. Für den L2 ist kein Triggerpuls, sondern nur eine *L2-accept-Nachricht* (L2A) bzw. *L2-reject-Nachricht* (L2R) vorgesehen, die spätestens $500 \mu\text{s}$ nach der Kollision am Subdetektor angekommen sein muss (siehe auch Kap. 3.3.4). Bei einem L2A werden die Daten vom Detektor abtransportiert und dem HLT zur Verfügung gestellt.

Als Veto stehen beim L2 vier past-future-protections zur Verfügung, die auch hier frei konfiguriert und gewählt werden können.

3.3.3. High-Level-Trigger

Ein weiteres Triggerlevel stellt der High Level Trigger (HLT) dar. Dazu wird mit einer Computer-Farm eine Online-Rekonstruktion der Daten vorgenommen und nach seltenen Ereignissen gesucht. Der HLT hat eine maximale Triggerrate von 30 Hz. Die interessanten Ereignisse werden dauerhaft mit dem DAQ-System gespeichert.

3.3.4. Trigger-Distributionsnetzwerk

ALICE benutzt wie andere Experimente am LHC das *Trigger, Timing and Control* (TTC) – System. Am entsprechenden Subdetektor befindet sich eine TTC-Partition, die sich aus drei Einheiten zusammensetzt: *Local Trigger Unit* (LTU), *TTC VMEbus Interface* (TTCvi) und *TTC Laser Transmitter* (TTCex). In Abb. 3.6 ist schematisch dargestellt, wie die TTC-Partition aufgebaut ist und mit dem CTP und der Detektorelektronik kommuniziert.

Das LTU stellt die Schnittstelle zwischen CTP und TTC-Partition dar. Es empfängt die Daten der CTP, die über das TTC-System verteilt werden. Der L0 wird direkt über ein differentielles LVDS-Kabel an die Detektorelektronik weitergegeben, um Verzögerungen zu minimieren⁵. Die restlichen Daten werden in zwei Kanälen an die Detektorelektronik übermittelt. Der A-Kanal ist dabei exklusiv für den L1-Puls reserviert und wird direkt an den Laser Transmitter TTCex weitergegeben. Über den B-Kanal werden Triggernachrichten (L1M, L2A, L2R) und verschiedene Befehle geschickt. Das

⁵ Dies trifft nicht auf den TRD zu (vgl. Kap. 4.5)

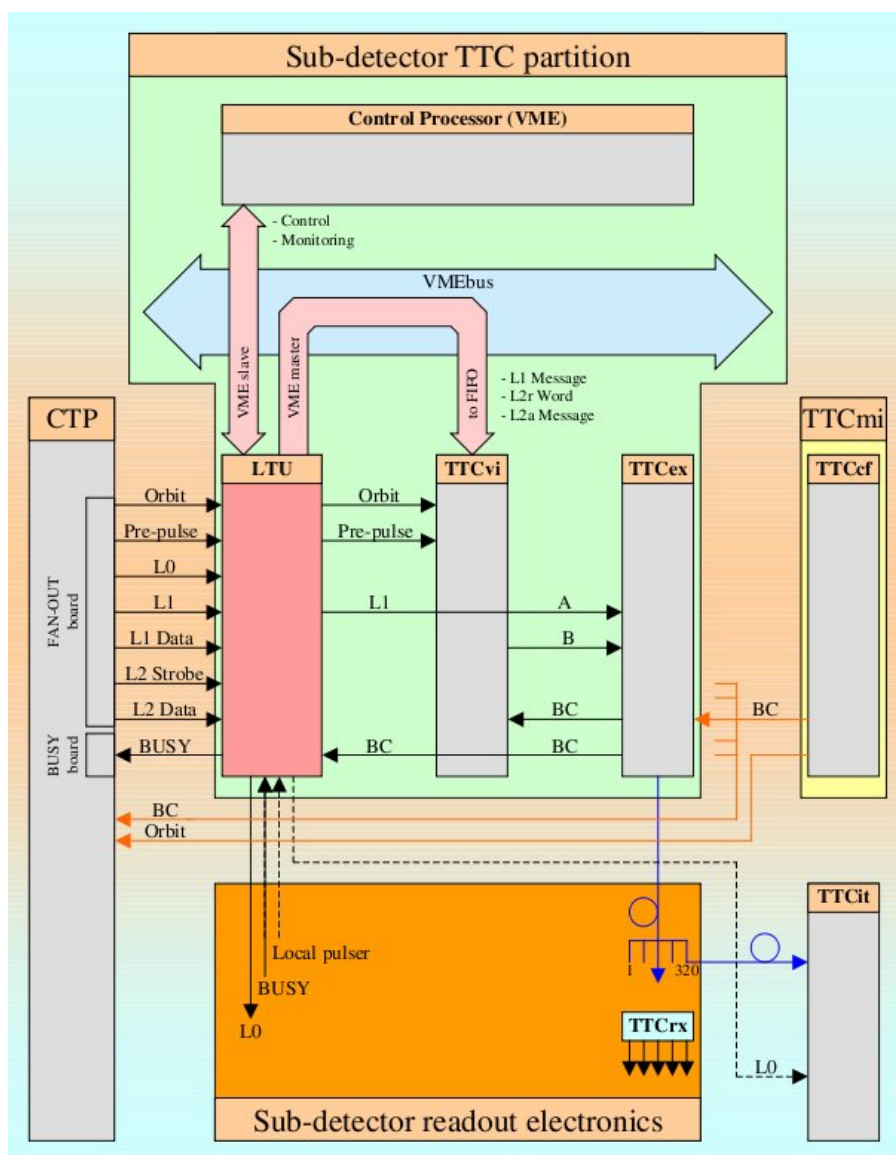


Abb. 3.6.: Das Trigger-System von ALICE. Die CTP empfängt Triggerbeiträge bzw. verteilt Trigger an die einzelnen Subdetektoren, wo sie jeweils von einer TTC-Partition verarbeitet und über das TTC-System übertragen werden. Über TTCmi wird den einzelnen Subdetektoren die TTC-CLK zur Verfügung gestellt. Bild aus [Jov04]

LTU empfängt die Nachrichten, die für den B-Kanal bestimmt sind, bringt sie in das TTC-Format und leitet sie über den VME-Bus an den TTCvi weiter.

TTCvi gibt den B-Kanal an den TTCex. Hier werden A- und B-Kanal gemultiplext, in ein optisches Signal gewandelt und an die Detektorelektronik über Glasfasern verteilt⁶.

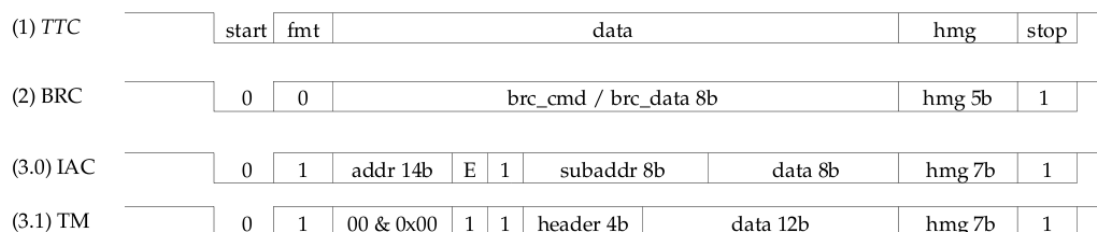
⁶ Dazu wird in der Regel noch ein TTCoc verwendet, der einen der zehn TTCex-Ausgänge auf bis zu 32 Fasern aufteilt.

Über das TTCmi wird die *bunch-crossing clock* (TTC-CLK), ein 40,079 MHz Taktgeber⁷ und die *orbit clock*, ein 11,426kHz Taktgeber, verteilt. Das sind die globalen Taktgeber des LHC, die vom Prevessin Kontrollraum verteilt werden. Zur Detektorelektronik wird die TTC-CLK nicht direkt übertragen. Sie muss dort von speziellen TTCrx-Chips wiederhergestellt werden, die außerdem Kanal A und Kanal B demultiplexen und in elektronische Signale wandeln.

Als Rückgabe von den Subdetektoren erhält der CTP ein *busy*-Signal und die jeweiligen Triggerbeiträge der Subdetektoren.

Triggernachrichten und Befehle über den B-Kanal

Es wird zwischen *Broadcast Commands* (BRC) und *Individually Addressed Commands* (IAC) unterschieden. BRCs sind an alle TTCrx im System adressiert und werden von ihnen ausgeführt. IACs sind an spezielle Empfänger adressiert und werden entweder vom TTCrx oder von der Elektronik, mit der er verbunden ist, ausgeführt. Wie das IAC verarbeitet werden soll, wird durch das *address space selection bit* (E) in der Nachricht bestimmt. In Abb. 3.7 ist das TTC-Format gezeigt und erklärt. Die Triggernachrichten L1M, L2A und L2R sind spezielle Formate von IACs, die an die



(b) Channel B.

- (1) General TTC frame format,
- (2) Broadcast frame (BRC),
- (3.0) Individually addressed frame (IAC) and
- (3.1) Trigger message, a special form of an IAC.

Abb. 3.7.: Das TTC-Datenformat für den B-Kanal. Oben ist das allgemeine Format gezeigt, in das sowohl BRC als auch IAC Befehle gebracht werden. Am Anfang und Ende sind immer Start- und Stop-Bit vorhanden, die einer logischen 1 bzw. einer logischen 0 entsprechen. Auf das Start-Bit folgt das FMT-Bit, das angibt, ob es sich um ein BRC oder IAC handelt. Bei BRC folgt darauf das 8-Bit Datenwort und fünf redundante Bits, die zur Fehlererkennung und -korrektur hinzugefügt sind. Bei IAC folgt auf das Start-Bit erst die Adresse des TTCrx Empfängers und adress space selection bit, das indiziert, ob der Befehl weitergeleitet werden soll. Falls dies der Fall ist wird die darauf folgende Subadresse und das 8-Bit Datenwort extern zur Verfügung gestellt. Triggernachrichten sind ein spezielles Format der IACs. Bild aus [Kir07]

Adresse 0 gesendet werden. Das bedeutet, dass sie wie BRCs an alle TTCrx gesendet, aber nicht von ihnen verarbeitet, sondern an die Elektronik weitergegeben werden. Außerdem hat der Befehl

⁷ Ein Takt der TTC-CLK entspricht der Zeit, die zwischen zwei passierenden Bunches vergeht. Bei stabilen Strahlen im LHC ist sie mit den Bunches synchronisiert.

nicht ein 8-Bit-, sondern ein 16-Bit-Datenwort, das sich in einen 4-Bit Kopf und 12-Bit für den eigentlichen Befehl aufteilt. Der Kopf gibt an, um welchen Typ Triggernachricht es sich handelt. Ein L1A besteht aus fünf Datenworten, in denen neben ereignisspezifischen Informationen, der Zustand der 50 Trigger-Klassen der CTP übermittelt wird. Ein L2A beinhaltet neben den Zuständen der Triggerklassen noch Informationen zur Identifizierung des Ereignisses, zu dem der L2A gehört, da in der Zeit zwischen L1A und L2A bereits neue Triggersequenzen gestartet werden dürfen. Er besteht aus acht Worten.

In diesem Kapitel wurden das Funktionsprinzip und die Eigenschaften des LHC diskutiert und anschließend das Experiment ALICE vorgestellt. Mit ALICE sollen die Eigenschaften eines Quark-Gluon-Plasmas, das bei Schwerionenkollisionen erzeugt wird, untersucht werden. ALICE besteht aus einzelnen Subdetektoren, deren Funktion in diesem Kapitel beschrieben wurde.

Für ALICE wurde außerdem ein komplexes Triggersystem entwickelt, das aus verschiedenen Levels besteht, da die Subdetektoren verschieden schnell Ergebnisse liefern und ausgelesen werden können.

4. Transition Radiation Detector

In diesem Kapitel wird der Transition Radiation Detector von ALICE diskutiert. Nach einer physikalischen Motivation für den TRD wird zunächst das Prinzip der Übergangsstrahlung (TR kurz für *Transition Radiation*) kurz erläutert. Anschließend wird der Aufbau und die Funktionsweise des TRD vorgestellt. Danach wird das Triggersystem von ALICE wieder aufgegriffen und speziell für den TRD erläutert.

Im Anschluss daran wird die Detektorelektronik diskutiert, mit der die Daten ausgelesen und verarbeitet werden. Anhand einer Triggersequenz wird der zeitliche Ablauf der Datenauslese erläutert.

4.1. Motivation für den TRD

Um ein QGP nachzuweisen und zu untersuchen, wird nach typischen Signaturen gesucht, die in Kap. 2.2.2 bereits diskutiert wurden. Bei einigen dieser Signaturen ist es nötig, Elektronen oder Positronen mit hohem Transversalimpuls $p_t \geq 1 \text{ GeV } c^{-1}$ zu identifizieren:

- J/Ψ und Υ können in Elektron–Positron–Paare mit hohem Impuls zerfallen.
- D– und B–Mesonen können über den semileptonischen Zerfall Elektronen oder Positronen mit hohem Impuls produzieren. Zum Beispiel D^+ zerfallen mit einer Wahrscheinlichkeit von ca. 9% zu [PDG10]:

$$D^+ \rightarrow \bar{K}_0 + e^+ + \mu_e \quad (4.1)$$

Problematisch beim Nachweis der erzeugten Elektronen ist, dass im relevanten Impulsbereich ($p_t \geq 1 \text{ GeV}$) Pionen die Teilchen sind, die am häufigsten in Schwerionenkollisionen produziert werden. π^- und π^+ tragen, genau wie Elektronen bzw. Positronen, eine elektrische Ladung von $\pm e$. Da sie nicht verwechselt werden dürfen, ist eine Unterdrückung der Pionen wichtig. Während die TPC eine ausreichende Unterdrückung des Pionenhintergrundes bei Transversalimpulsen bis $1 \text{ GeV } c^{-1}$ leistet, ist sie für höhere Impulse nicht geeignet. Um Pionen mit höherem Impuls zu unterdrücken und die Elektronen dabei zu identifizieren, wurde der TRD zu ALICE hinzugefügt.

In einem TRD–Modul kann Übergangsstrahlung von geladenen Teilchen erzeugt und anschließend nachgewiesen werden. Im relevanten Impulsbereich wird Übergangsstrahlung nur von Elektronen erzeugt, was eine effektive Pionenunterdrückung ermöglicht (siehe Kap. 4.4).

Mit dem TRD können Elektronen bzw. Positronen mit hohem Impuls nicht nur in der Analyse identifiziert werden. Der TRD ist außerdem in der Lage, über verschiedene Methoden einen schnellen Trigger-Beitrag (L1) zu liefern, der physikalisch interessante Ereignisse markiert, damit sie bevorzugt aufgezeichnet werden. Dieser Trigger ist zum Beispiel wichtig, um zu gewährleisten, dass eine ausreichende Statistik für den Zerfall von Υ in Elektron-Positron-Paare erreicht wird, da während eines typischen Betriebsjahres mit Blei-Blei-Kollisionen nur ca. 550 solche Zerfälle bei ALICE erwartet werden [ALI01].

In den Driftkammern der TRD-Module wird nicht nur die Übergangsstrahlung, sondern auch die Spur der Teilchen aufgezeichnet. Der TRD trägt damit zur globalen Spurrekonstruktion von ALICE bei. Über die Krümmung der Spur durch das Magnetfeld des L3-Magneten kann der Impuls von geladenen Teilchen bestimmt werden.

4.2. Übergangsstrahlung

Um Elektronen mit hohem Impuls zu identifizieren, wird Übergangsstrahlung nachgewiesen, die entstehen kann, wenn ein geladenes Teilchen eine Grenzfläche zweier Medien mit verschiedenen Dielektrizitätskonstanten passiert. Die Wahrscheinlichkeit, dass Übergangsstrahlung erzeugt wird, hängt vom Lorentz-Faktor $\gamma = \frac{E}{mc^2}$ ab. Erst ab einer Schwelle von $\gamma \gtrsim 1000$ [Kwe09] wird Übergangsstrahlung erzeugt.

Für Elektron mit einem Transversalimpuls $p_t = 3 \text{ GeV } c^{-1}$ beträgt der Lorentz-Faktor $\gamma_e \approx 6000$. Ein Pion ist mit $m_\pi \approx 140 \text{ MeV } c^{-2}$ wesentlich schwerer als ein Elektron ($m_e \approx 511 \text{ keV } c^{-2}$) und hat bei gleichem Impuls einen Lorentz-Faktor $\gamma_\pi \approx 20$. Während Elektronen im relevanten Impulsbereich ($p_t \geq 1 \text{ GeV}$) Übergangsstrahlung erzeugen, liegt der Lorentz-Faktor von Pionen somit unter der Schwelle und sie erzeugen keine Übergangsstrahlung.

Die Wahrscheinlichkeit, dass Übergangsstrahlung bei einem einzelnen Materialübergang erzeugt wird, ist trotz der hohen Lorentz-Faktoren der Elektronen relativ gering. Daher ist es wichtig, möglichst viele Materialübergänge im Radiator zu haben, damit die Wahrscheinlichkeit steigt und jedes Elektron mindestens ein Photon produziert.

Der Emissionswinkel φ der Übergangsstrahlung relativ zur Spur des Teilchens ist proportional zu γ^{-1} [ALI01]. Übergangsstrahlung wird also in einem engen Kegel um die Trajektorie des Teilchens emittiert. Im Gasdetektor des TRD ionisiert sowohl die Übergangsstrahlung als auch das Teilchen entlang seiner Trajektorie das Zählgas. Da der Winkel der Übergangsstrahlung zur Trajektorie des Teilchens klein ist, kann die Übergangsstrahlung der Spur des Teilchens zugeordnet werden, das die Übergangsstrahlung erzeugt hat. (vgl. Abb. 4.3).

4.3. Aufbau des TRD

In Abb. 4.1 ist der Aufbau des TRD skizziert. Er ist in azimuthaler Richtung in achtzehn 8 m lange Segmente, die sogenannten Supermodule, unterteilt und deckt damit den kompletten azimuthalen

Bereich und eine Pseudo-Rapidity von $-0,9 < \eta < 0,9$ ab. Die Supermodule sind in Strahlrichtung in jeweils 5 Stapel (*stacks*) und in radialer Richtung in 6 Lagen (*layer*) unterteilt. Der gesamte TRD setzt sich somit aus 540 Modulen zusammen. Auf jedem Modul sind je nach Größe 96 oder 128 *Muti-Chip-Module* (MCM)¹ angebracht, die jeweils 18 Kanäle auslesen und parallel eine erste Auswertung der Daten vornehmen, um einen schnellen Trigger zu ermöglichen.

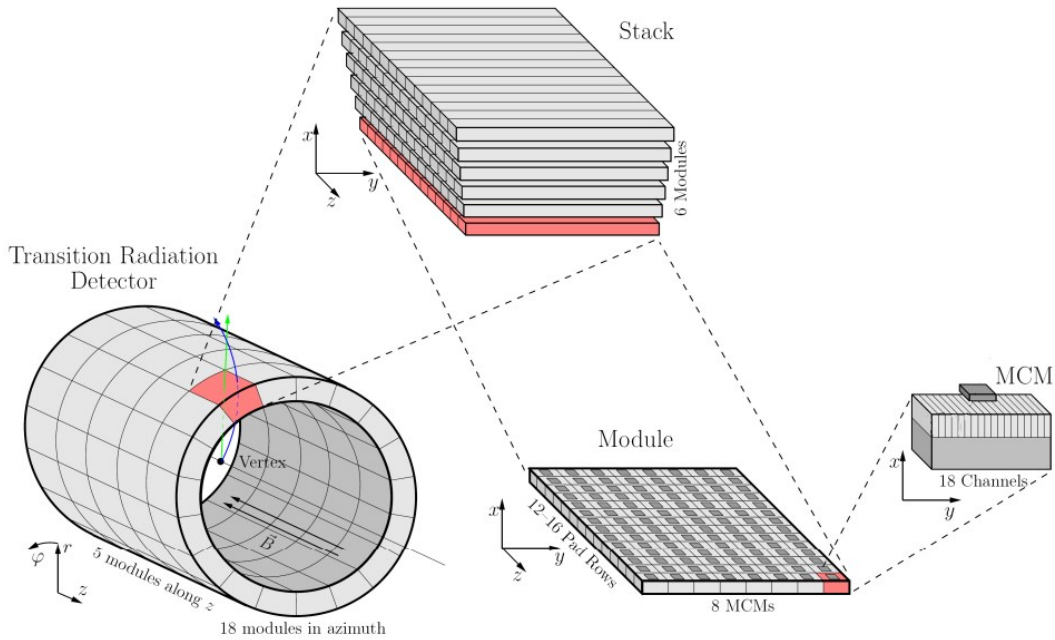


Abb. 4.1.: Der TRD und seine Zusammensetzung aus den einzelnen Supermodulen bzw. Modulen. Jedes Supermodul ist in fünf Stacks mit sechs Lagen unterteilt und beinhaltet somit 30 unabhängige Module. Auf jedem Modul sind 96 oder 128 MCMs zur Auslese und ersten Auswertung der Daten angebracht. Bild aus [Cuv03]

4.3.1. Modulaufbau

In Abb. 4.2 ist ein Querschnitt durch ein Modul gezeigt. Es ist aus verschiedenen Schichten aufgebaut. Diese sind in radialer Richtung: Radiator, Auslesekommer und Trägerplatte (*backplane*). Die Auslesekommer unterteilt sich in einen Drift- und einen Verstärkungsbereich.

Radiator

In dem Radiator wird die Übergangsstrahlung erzeugt. Üblicherweise werden oft Stapel mit mehreren hundert Folien aus Polypropylen als Radiator benutzt, um Übergangsstrahlung zu erzeugen, da sie die beste Ausbeute bieten [ALI01]. An den Radiator des TRD sind jedoch zusätzliche Anforderungen gestellt:

¹ Diese werden in Kap. 4.6 noch diskutiert.

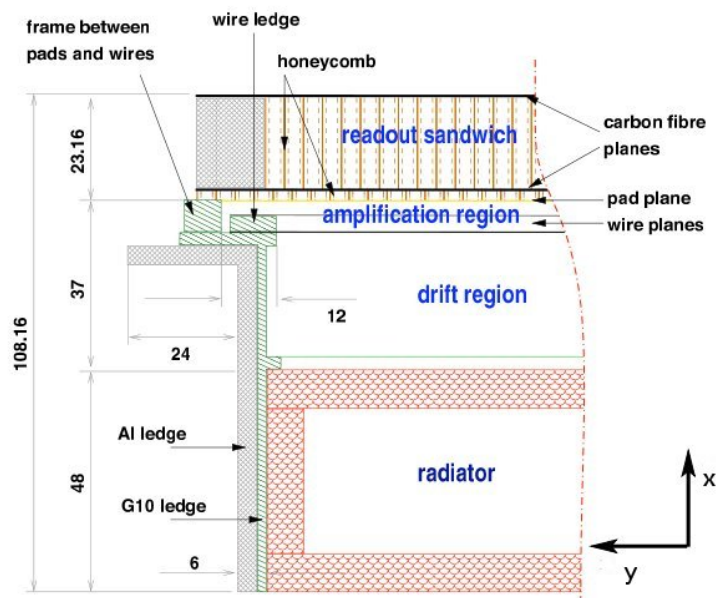


Abb. 4.2.: Ein Querschnitt durch ein TRD-Modul. Im unteren Bereich befindet sich der Radiator. Daran schließt sich in radialer Richtung der Driftbereich, der Verstärkungsbereich, die Auslese-Pads und die *backplane* an. Mit *wire planes* sind hier die zwei Drahtebenen beschriftet. Alle Längen sind in Millimeter angegeben. Bild aus [Ems09]

- Gute TR-Ausbeute, bei minimaler Dicke und Dichte.
- Hohe mechanische Stabilität, damit kein totes Material zur Stabilisierung in den L3-Magneten eingebracht werden muss, da dies zu Akzeptanzverlusten führen würde.

Der Radiator für den TRD besteht aus einer Schichtstruktur aus Rohacell HF71 Schaum und Polypropylen-Fasern mit einer Gesamtdicke von 4,8 cm. Die Schichten aus Rohacell an den Grenzflächen des Radiators gewährleisten eine hohe mechanische Stabilität, bei geringer Dichte und relativ guter TR-Ausbeute. Dazwischen befinden sich Polypropylen-Fasern, die als Haupt-Radiatormaterial eine ähnliche TR-Ausbeute wie Folienstapel aus Polypropylen liefern [ALI01]. Zusätzlich ist auf jeder Seite des Moduls eine dünne Glasfaserschicht zur Verstärkung angebracht.

Mit diesem Radiator werden von einem Elektron mit $\gamma \geq 1000$ durchschnittlich 1,45 Röntgenphotonen im Bereich von 1 keV bis 30 keV erzeugt [ALI08a].

Auslesekommer

Zur Detektion der Übergangsstrahlung und der Trajektorien der Teilchen wird eine Driftkammer verwendet. Direkt auf den Radiator ist eine mit Aluminium beschichtete Mylar-Folie geklebt, die als Eintrittsfenster in die Driftregion und als Driftelektrode dient. Die Driftregion ist 3 cm lang und endet mit den Kathodendrähten, die in einem Abstand von 2,5 mm angeordnet sind. Hinter der

Ebene der Kathodendrähte befinden sich die Anodendrähte. Sie sind in der Mitte der 7 mm langen Verstärkungsregion in einem Abstand von 5 mm angebracht.

Die Verstärkungsregion endet mit den Kathodenpads. Das sind im Durchschnitt $6,3 \text{ cm}^2$ große Streifen, die auf der Ausleseplatte angebracht sind. Die Kathodenpads sind relativ zur z-Achse um $\alpha = 2^\circ$ gedreht. In zwei aufeinander folgende Lagen im Supermodul sind sie in entgegengesetzte Richtung gedreht. Durch diese Drehungen wird die Auflösung in z-Richtung erhöht.

Die *backplane* besteht aus einer Wabenstruktur und dient lediglich der mechanischen Stabilität. Auf der äußeren Seite sind die MCMs angebracht, die jeweils 18 Kathodenpads auslesen (vgl. Kap. 4.6).

Das Innere der Kammer, also die Drift- und Verstärkungsregion, ist mit einem Xe-CO₂-Gemisch gefüllt.

4.4. Funktionsprinzip

In Abb. 4.3 ist das Funktionsprinzip einer Auslesekammer schematisch dargestellt. Passiert ein hochenergetisches geladenes Teilchen die Driftregion, ionisiert es entlang seiner Spur die Xenon-Atome. Ein Elektron erzeugt zusätzlich noch Übergangsstrahlung im Radiator, die ebenfalls das Xenon-Gas ionisiert.

Spur der Teilchen

An die Driftelektrode wird ein Potenzial von $-2,1 \text{ kV}$ angelegt, wodurch ein homogenes elektrisches Feld zwischen dieser und den Kathodendrähten erzeugt wird. Die herausgelösten Elektronen werden durch das Feld beschleunigt und driften zu den Kathodendrähten. Durch die konstante Driftgeschwindigkeit ist die Driftzeit proportional zur radialen Position, an der das Elektron herausgelöst wurde. Elektronen die direkt nach dem Radiator erzeugt werden, brauchen dazu länger als solche die näher an den Kathodenpads erzeugt wurden. Bei einer Spannung von $-2,1 \text{ kV}$ beträgt die maximale Driftzeit $2 \mu\text{s}$.

In Kap. 2.3.1 wurde bereits der Energieverlust von geladenen Teilchen in Materie diskutiert. Durch die Ionisation der Xe-Atome deponieren geladene Teilchen eine spezifische Energie pro Strecke $\frac{dE}{dx}$ in der Kammer. Während Pionen im relevanten Impulsbereich noch minimal ionisierende Teilchen sind oder sich im relativistischen Anstieg befinden, liegen Elektronen im Fermi-Plateau. Dadurch deponieren Elektronen bei gleichem Impuls mehr Energie in der Kammer als Pionen.

Übergangsstrahlung

Zusätzlich zu der Ionisationsspur erzeugen Elektronen mit $\gamma = 1000$ durchschnittlich 1,45 TR-Photonen im Radiator. Durch die geringe mittlere freie Weglänge der TR-Photonen in Xenon werden sie mit hoher Wahrscheinlichkeit am Anfang der Kammer durch Photo- und Compton-Effekt gestoppt. Ein typisches Übergangsstrahlungsphoton mit $E = 10 \text{ keV}$ hat eine mittlere freie Weglänge

von 1 cm [ALI01]. So wird ein räumlich begrenztes Elektronensignal erzeugt. Da der Abstrahlungswinkel der Übergangsstrahlung zur Flugbahn des Elektrons invers proportional zum Lorentz-Faktor ist, überlagert sich das Signal der Übergangsstrahlung mit dem der Ionisationsspur des Teilchens.

Verstärkung

Die Ladung, die durch das ionisierende Teilchen und die Übergangsstrahlung in der Kammer erzeugt wird, ist nicht ausreichend, um direkt ausgelesen zu werden. Deshalb muss sie im Verstärkungsbe- reich vervielfacht werden. An den Anodendrähten liegt eine Spannung von +1,4 kV an. Dadurch wird ein starkes inhomogenes Feld erzeugt, das die Elektronen lawinenartig um einen Faktor von ca. 5000 [ALI04] verstärkt. Die erzeugten Elektronen werden durch die Anodendrähte abgesaugt, während die Ionen langsam zu den Kathodenpads driften.

Die ladungssensitiven Vorverstärker in der Ausleseelektronik wandeln die auf den Kathodenpads induzierte Ladung in ein Spannungssignal, das dann weiterverarbeitet wird. Zurück bleiben angeregte Xenon-Atome, die durch Emission von Photonen wieder in den Grundzustand gelangen. Damit diese Photonen die Lawine nicht dauerhaft aufrecht erhalten und eine neue Spur aufgenommen werden kann, wird ein Löschgas (*quench gas*) zum eigentlichen Zählgas hinzugefügt. Im TRD werden dem Xenon dazu 15% CO₂ beigemischt, das die Photonen, die das Xenon emittiert, absorbiert.

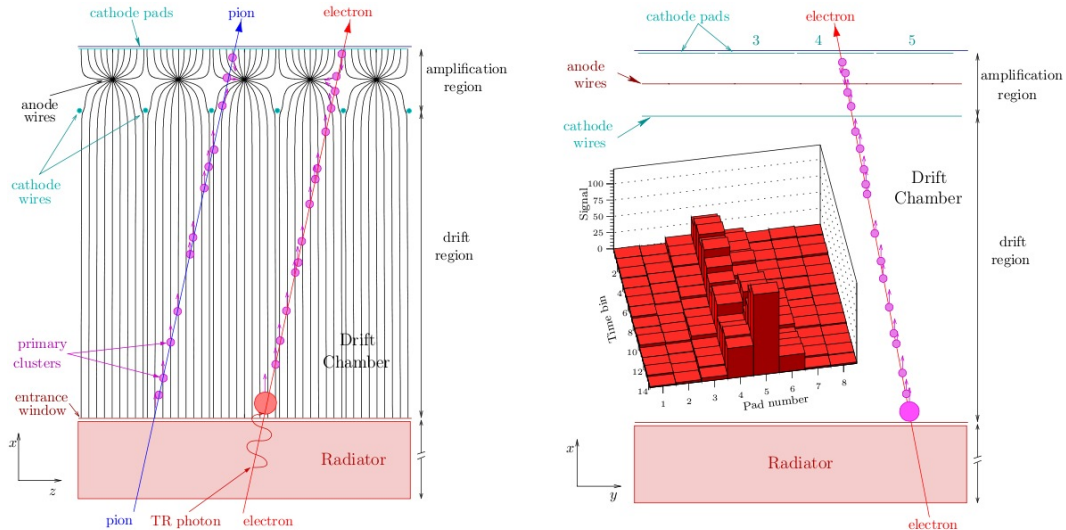


Abb. 4.3.: Die Funktionsweise des TRD ist schematisch gezeigt. Links ist ein Schnitt durch eine Kammer gezeigt und an welcher Position Elektronen bzw. Pionen ein Signal erzeugen. Rechts ist die induzierte Ladung auf den Kathodenpads pro Driftzeit aufgetragen. Elektronen erzeugen zusätzlich zur Ionisationsspur noch Übergangsstrahlung. Diese wird am Anfang des Driftvolumens gestoppt und generiert eine Signalmaximum bei hohen Driftzeiten. Bild aus [ALI01]

Signalintepretation

In Abb. 4.4 ist die erzeugte Ladung in Abhängigkeit der Driftzeit für Elektronen mit und ohne Radiator und für Pionen mit Radiator gezeigt. Sowohl Elektron als auch Pion haben in diesem Diagramm einen Impuls von 2 GeV.

Das Maximum bei hohen Driftzeiten bei Elektronen ist der Übergangsstrahlung zuzuordnen. Außerdem ist deutlich zu erkennen, dass der spezifische Energieverlust von Elektronen höher ist als der von Pionen.

Das Maximum bei niedrigen Driftzeiten beruht darauf, dass die Driftzeit im Verstärkungsbereich für Elektronen, die mit gleichem Abstand vor und hinter den Anodendrähten erzeugt werden, gleich ist. Es wird deshalb auch Verstärkungsmaximum genannt und ist unabhängig vom Teilchentyp.

Aus diesen typischen Merkmalen im generierten Signal gehen einige Methoden hervor, wie Elektronen von Pionen unterschieden werden können. Die einfachste Methode ist, die gesamte deponierte Ladung Q eines Teilchens mit einer Wahrscheinlichkeitsverteilung (*Likelihood*) zu vergleichen (LQ-Methode). Die Wahrscheinlichkeitsverteilung gibt an, wie wahrscheinlich eine deponierte Ladung für Elektronen bzw. Pionen ist. Eine verbesserte Pionenunterdrückung wird jedoch erreicht, wenn zusätzlich noch der Ort der maximalen Primärladung berücksichtigt wird. Dadurch können die großen Ladungcluster der Übergangsstrahlung identifiziert werden. Eine weitere vielversprechende Methode ist es, neuronale Netzwerke zu trainieren (Details zur Elektron-Pion-Separation sind in [Wil04] und [Wil10] zu finden). Für Impulse $> 3 \text{ GeV } c^{-1}$ wird bei zentralen Pb-Pb-Kollisionen eine Pion-Unterdrückung von 100 bei einer Elektroneneffizienz von 90 % angestrebt [ALI01]. Das bedeutet, dass 90 % aller Elektronen korrekt als solche identifiziert werden, während 1% der Pionen fälschlicherweise als Elektronen identifiziert werden.

Zusätzlich wird mit den TRD noch die Spur und der Impuls der Teilchen bestimmt. Die Ortsauflösung in r - φ liegt bei $400 \mu\text{m}$ und bei 2 mm in z -Richtung. Es wird eine Impulsauflösung von 5 % bei Impulsen von $5 \text{ GeV } c^{-1}$ erreicht [ALI01].

4.5. Triggersystem beim TRD

In Kapitel 3.3 wurde das Triggersystem von ALICE bereits vorgestellt. Für den TRD ergeben sich insbesondere durch den benötigten Pretrigger einige Besonderheiten, die in diesem Kapitel erläutert werden sollen.

Auf den Modulen des TRD ist massiv parallele Hardware implementiert, um eine Vorabauswertung der Daten zu ermöglichen, mit der ein schneller L1-Beitrag an die CTP geschickt werden kann. Um Energie zu sparen, wird die Elektronik die meiste Zeit in einem Standby-Modus betrieben, aus dem sie durch den Pretrigger aufgeweckt wird.

Da der TRD vor dem Pretrigger keine Daten aufzeichnet, ist die Latenz kritisch. Aus diesem Grund befindet sich das Pretrigger-System in dem L3-Magneten und der PT wird exklusiv für den TRD generiert. Dazu werden die Daten der Subdetektoren T0, V0 und TOF ausgewertet.

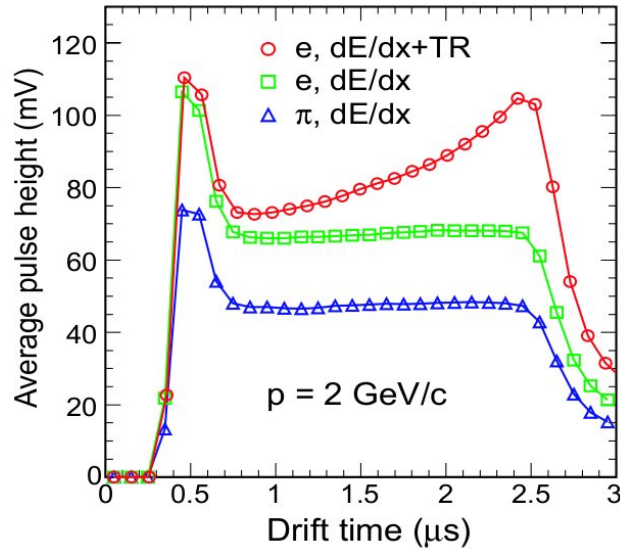


Abb. 4.4.: Die erzeugten Signale von Elektronen mit und ohne Radiator und von Pionen mit Radiator. Deutlich ist das Maximum bei langen Driftzeiten bei Elektronen zu erkennen, das durch die Übergangsstrahlung erzeugt wird. Bei kurzen Driftzeiten befindet sich das Verstärkungsmaximum, das unabhängig vom Teilchentyp beobachtet wird. Neben der Übergangsstrahlung, die durch Elektronen erzeugt wird, ist der spezifische Energieverlust durch Ionisation der Xe-Atome von Elektronen größer als der von Pionen. Bild aus [ALI08a]

Zusätzlich zur Pretrigger-Generation muss das Pretrigger-System das TTC-Signal empfangen und es zusammen mit dem Pretrigger an die 18 Supermodule verteilen.

4.5.1. Aufbau des TRD-Triggersystems

In Abb. 4.5 ist das Trigger-System für den TRD schematisch dargestellt. Es lässt sich unterteilen in einen oberen Teil, der sich außerhalb des L3-Magneten und einen unteren Teil, der sich innerhalb des Magneten befindet. Mit CTP/LTU ist hier die TTC-Partition, also die Schnittstelle zwischen Detektorelektronik und CTP, gekennzeichnet.

Zentrales Element ist die Kontrollbox B (CB-B kurz für *Control Box B*)². Dort laufen alle Signale zur Pretrigger-Generation und das TTC-Signal zusammen, werden ausgewertet und gegebenenfalls an die Supermodule verteilt.

Pretrigger-Generation

Der Pretrigger wird aus Triggerbeiträgen von T0, V0 und TOF generiert. V0 besteht aus jeweils 32 Photovervielfachern (PMT) auf der A- und C-Seite, deren Signale zu acht in einer *Front-End-Box* (FEB) gesammelt und verstärkt werden. T0 besteht aus zwölf PMTs auf der A- und zwölf auf der C-

² Die anderen beiden Kontrollboxen, die sich jeweils auf der A- bzw. C-Seite befinden, werden mit CB-A bzw. CB-C abgekürzt.

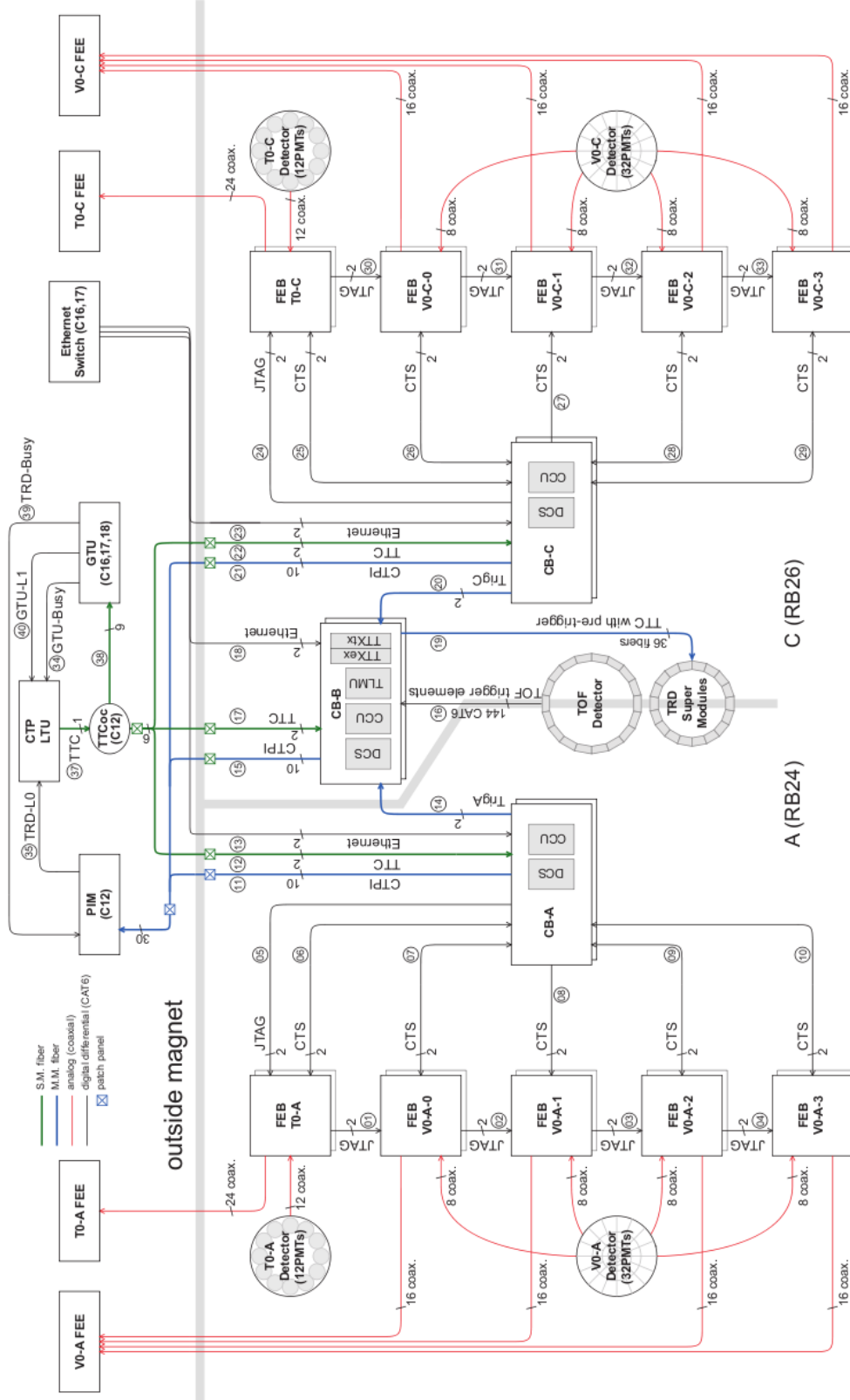


Abb. 4.5.: Das Triggersystem des TRD. V0 und T0 liefern Beiträge zur Pretrigger-Generation, die über FEBs und zwei Kontrollboxen gesammelt und an CB-B weitergegeben werden. In CB-B werden die Daten zusammen mit den Beiträgen von TOF ausgewertet und eine endgültige Pretrigger-Entscheidung getroffen. Außerdem empfängt CB-B das TTC-Signal. Mit einer Zustandsmaschine werden PT-, L0A- und L1A-Triggerpulse für die 18 Supermodule generiert und über optische Fasern übertragen. Die GTU gibt ein busy-Signal und einen L1-Beitrag an den CTP zurück. Bild aus [PTS10]

Seite, die auch jeweils in einer FEB zusammengefasst werden. Da V0 und T0 ihre eigene Elektronik außerhalb des L3-Magneten haben, um die Signale weiterzuverarbeiten, werden sie sowohl dorthin, als auch zu CB-A bzw. CB-C zur Pretrigger-Erzeugung gesendet. CB-A/C sammeln die Daten von je fünf FEBs und geben sie über optische Fasern an CB-B weiter. Die Daten von TOF werden direkt an CB-B gegeben. CB-B wertet die Pretrigger-Beiträge der drei Subdetektoren aus und trifft eine endgültige Entscheidung. Wie die einzelnen Triggerbeiträge sein müssen, damit ein Pretrigger generiert wird, ist flexibel einstellbar.

Für den Pretrigger ist, wie für die L0-L2-Trigger, eine past-future-protection vorhanden, die sicherstellt, dass sich Ereignisse nicht anstauen oder überschneiden.

Trigger-Distribution

Von der CTP empfängt das LTU das Standard-Signal, das L0, L1 und L2 beinhaltet. Über den A-Kanal des TTC-Signals wird am TRD nicht wie bei den anderen Subdetektoren nur der L1A, sondern auch der L0A gesendet. Das TTC-Signal wird über TTCex und TTCoc an die GTU und die CB-B verteilt. In der CB-B laufen alle Trigger-Signale zusammen und werden an die Supermodule verteilt.

Die Detektorelektronik auf den Supermodulen erwartet insgesamt drei Triggerpulse: PT, L0A und L1A³. Alle Triggerpulse sind genau einen TTC-Takt lang (25 ns) und werden über je zwei Fasern⁴ an die Supermodule weitergegeben. Damit die drei identischen Pulse als PT, L0A und L1A identifiziert werden können, dürfen sie nur in dieser Reihenfolge und in fest definierten Zeitfenstern relativ zum PT geschickt werden. In CB-B ist dazu eine Zustandsmaschine (FSM) eingebaut, mit der die Triggerpulse zur entsprechenden Zeit generiert werden, falls sie von der CTP geschickt werden. Sie ist schematisch in Abb. 4.6 dargestellt.

In der Detektorelektronik ist eine ähnliche FSM eingebaut, die gewährleistet, dass nur Trigger akzeptiert werden, die in ein gültiges Zeitfenster fallen. Sie wird in Kapitel 7.3.1 näher erläutert.

Außerhalb des L3-Magneten ist das *Pretrigger-Interface-Module* (PIM) platziert, das die Pretrigger-Daten von CB-A/B/C über optische Fasern sammelt und daraus einen L0-Beitrag für den CTP generiert.

4.6. Detektorelektronik

Ein großer Teil der Detektorelektronik des TRD ist direkt auf den Modulen verbaut. Auf den Modulen befinden sich jeweils sechs bis acht Ausleseplatinen (ROB kurz für *ReadOut Board*)⁵, auf denen *Multi-Chip-Module* (MCM), optische Ausleseseinheiten (ORI kurz für *Optical Readout*

³ Nur der L1-Puls und nicht die Nachricht

⁴ Beide Fasern übermitteln das gleiche Signal und es wird nur eine Faser benötigt. Die zweite Faser ist nur aus Gründen der Redundanz vorhanden.

⁵ Es gibt zwei verschiedene Modulgrößen. Die Module im mittleren Stapel sind kleiner als die anderen und haben dementsprechend zwei ROB weniger.

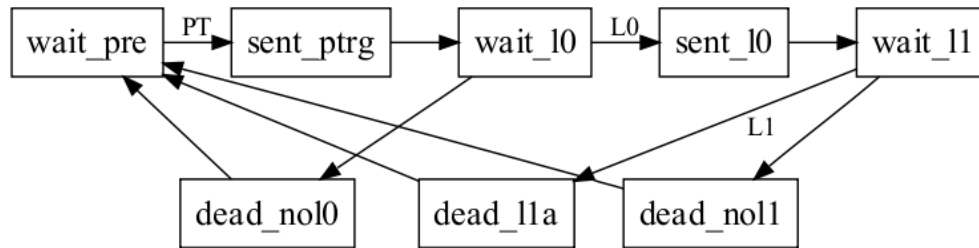


Abb. 4.6.: Zustandsmaschine in CB-B zur Trigger-Generation für die Supermodule. Das Supermodul erwartet einen PT und als Bestätigung einen L0A und L1A. In wait_pre wird auf einen Pretrigger gewartet und die Sequenz gestartet. L0A und L1A werden nur zu klar definierten Zeiten (in Bezug auf den Pretrigger) an die Supermodule ausgegeben. Die Sequenz wird abgebrochen, wenn kein L0A oder kein L1A von der CTP geschickt wird. Wird die Sequenz abgebrochen oder abgeschlossen, geht die FSM nach einer Wartezeit (dead_nol0, dead_nol1 und dead_l1a) zurück in wait_pre. Bild aus [Die10]

Interface) und Detektor-Kontrollkarten (DCS-Board kurz für *Detector Control System – Board*) angebracht sind. Außerhalb des L3-Magneten befindet sich die *Global Tracking Unit* (GTU).

In Abb. 4.7 ist die Detektorelektronik und in Abb. 4.8 der zeitliche Verlauf einer Triggersequenz schematisch dargestellt. In den MCMs werden die auf den Kathodenpads erzeugten Ladungen ausgelesen, verstärkt, gefiltert und in einem Zwischenspeicher abgelegt. Zusätzlich berechnen die MCMs Spursegmente (*tracklet*), also den Teil einer Spur, der sich innerhalb eines Moduls befindet, und geben diese über die ORIs an die GTU weiter. Die GTU fügt die einzelnen Spursegmente zu Spuren (*track*) zusammen, bestimmt den Impuls der Teilchen und sendet einen L1-Triggerbeitrag an den CTP. Wird von der CTP ein L1A zurückgegeben, werden die zwischengespeicherten Rohdaten an die GTU gesendet. Dort werden sie erneut bis zur L2-Entscheidung des CTP zwischengespeichert. Die einzelnen Komponenten dieser Auslekette werden im Folgenden erklärt.

4.6.1. Detector Control System

Der TRD wird über das Detector Control System konfiguriert und kontrolliert. Dazu ist auf jedem TRD-Modul ein DCS-Board angebracht, auf denen ein Linux-System installiert ist. Sie sind über Netzkabel in das DCS-Netzwerk eingebunden und können dadurch über ein TCP/IP-Netzwerk angesprochen und gesteuert werden.

Da die MCMs keinen dauerhaften Speicher haben, in dem sie Ihre Konfiguration ablegen können, wenn sie ausgeschaltet werden, werden sie über den *Slow Control Serial Network-Bus* (SCSN) vom DCS-Board konfiguriert.

Auf dem DCS-Board ist außerdem ein TTCrx vorhanden, der die TTC-CLK wiederherstellt und zusammen mit den Triggern von CB-B (PT, L0 und L1) an die MCMs weitergibt.

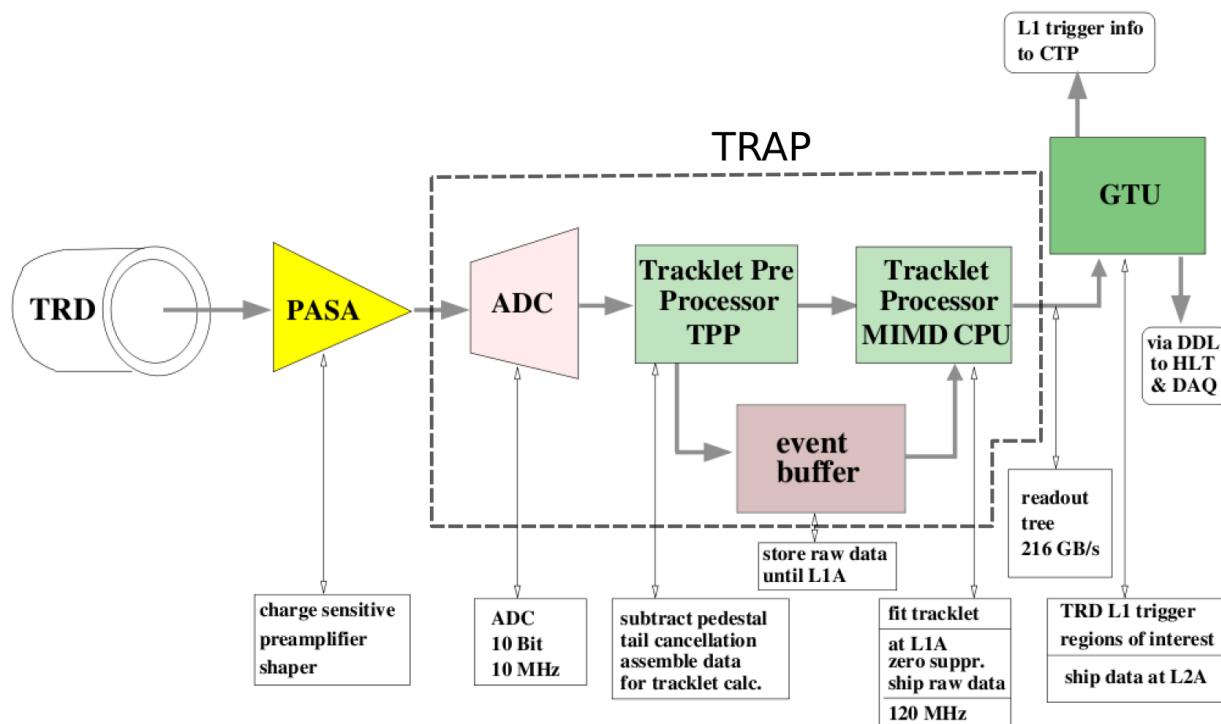


Abb. 4.7.: Schematische Darstellung der Signalverarbeitung der Detektorelektronik. Bis auf die GTU befinden sich alle Komponenten direkt auf den Modulen. Zunächst werden die Kathodenpads ausgelesen und mit dem PASA verstärkt. Nachdem die analogen Daten in digitale gewandelt wurden, werden digitale Filter angewandt. Der TPP bereitet die Spursegmentberechnung für den TP vor, während die Rohdaten in einem Zwischenspeicher abgelegt werden. Der TP berechnet bis zu vier Spursegmente, die über die optischen Verbindungen zur GTU weitergeleitet und dort ausgewertet werden. Die GTU liefert einen L1-Triggerbeitrag und speichert die Rohdaten nach einem L1A von der CTP. Nach einem L2A werden die Daten an das DAQ-System und den HLT weitergegeben. Bild aus [ALI01]

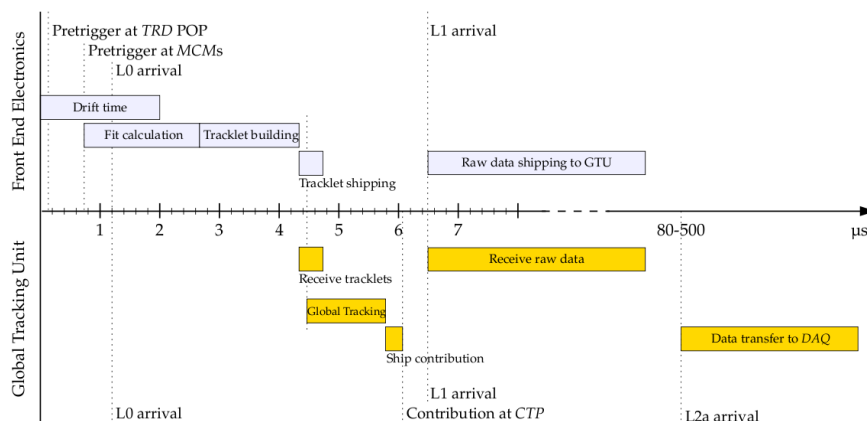


Abb. 4.8.: Zeitlicher Verlauf einer Triggersequenz bis zum L2A mit den einzelnen Arbeitsschritten der Detektorelektronik. Bild aus [Kir07]

4.6.2. Multi-Chip-Modul

Auf jedem ROB sind 16 MCMs angebracht, die jeweils 18 Kathodenpads auslesen. Zusätzlich gibt es auf jedem ROB einen oder zwei MCMs, die die Daten mehrerer MCMs für die optische Übertragung zur GTU zusammenfassen.

Ein MCM besteht aus einem *PreAmplifier/ShAper* (PASA) und einem *TRAcklet Processor* (TRAP). Der PASA ist ein Vorverstärker mit 18 Eingangskanälen und 21 Ausgangskanälen. Pro Kanal wird jeweils die induzierte Ladung von einem Kathodenpad in eine Spannung umgewandelt und an den TRAP weitergegeben. Über die drei zusätzlichen Ausgänge werden ein oder zwei Pad-Signale an die TRAPs benachbarter MCMs gegeben, um auch an den MCM-Grenzen eine Spurrekonstruktion zu ermöglichen.

Der TRAP kombiniert mehrere Aufgaben und ist über Register konfigurierbar, um z.B. festzulegen, welche Filter verwendet werden oder um diverse Verzögerungen anzupassen (für detaillierte Informationen sei auf [Ang05] verwiesen). Zunächst digitalisiert der TRAP die Eingänge mit 21 unabhängigen Analog-Digital-Wandlern (ADC). Sie arbeiten mit einer Frequenz von 10 MHz. Das entspricht 20 Datenpunkten (*timebins*) für eine $2\text{ }\mu\text{s}$ lange Driftzeit. Da der Pretrigger eine Laufzeit von einigen 100 ns hat, wird die Ausgabe der ADCs verzögert. Die Verzögerung lässt sich über ein Konfigurationsregister zwischen 100 und 600 ns einstellen. Typischerweise werden 30 Timebins aufgenommen, so dass vor und nach der eigentlichen Ionisationsspur noch einige Datenpunkte vorhanden sind.

Danach werden verschiedene digitale Filter angewandt, um das Signal aufzubereiten. Die einzelnen Filter sind je nach Konfiguration der TRAPs aktiviert bzw. deaktiviert. Sind sie aktiviert, dann laufen sie ständig. Ihre Ausgabe wird bei Ankunft eines Pretriggers an die Spur-Präprozessoren (TPP kurz für Tracklet PreProcessor) weitergegeben und gleichzeitig in einem Ereignispuffer gespeichert.

Die TPPs finden Treffer auf den Kathodenpads und berechnen Fit-Parameter für die Spursegmente. Sie arbeiten parallel zu den ADCs und beenden ihre Arbeit fast gleichzeitig mit der Driftzeit⁶. Die Fit-Parameter werden vom Spur-Prozessor (TP kurz für *Tracklet Processor*) ausgelesen und weiter verarbeitet.

Im TP arbeiten vier CPUs gleichzeitig, die jeweils ein Spursegment berechnen können. Die Spur der Teilchen ist durch den L3-Magneten gekrümmt. Da die Krümmung im interessanten Impulsbereich auf der kurzen Driftstrecke von 3 cm jedoch zu vernachlässigen ist, werden lineare Fits für die Spursegmente durchgeführt. Die Spursegmente sind mit nur einem 32-Bit Datenwort sehr klein im Vergleich zu den Rohdaten und werden $4,4\text{ }\mu\text{s}$ nach der Kollision über die ORIs an die GTU gesendet.

Sollte der CTP einen L1-Trigger schicken, werden die Rohdaten zur GTU übertragen. Sobald die Übertragung fertig ist, gehen die TRAPs wieder in den Ruhezustand und sind bereit für den nächsten Pretrigger.

⁶ Ist eine Verzögerung bei der Ausgabe der ADCs eingestellt, verschiebt sich der Zeitpunkt, an dem die TPPs fertig sind, um diesen Wert.

4.6.3. Global Tracking Unit

Die GTU sammelt die Daten der MCMs, liefert einen L1-Beitrag für den CTP und speichert die Rohdaten zwischen. Es besteht aus 90 *Track Matching Units* (TMU), 18 *SuperModule Units* (SMU) und eine *Trigger Generation Unit* (TGU). Der schematische Aufbau ist in Abb. 4.9 gezeigt. TMU, SMU und TGU basieren auf großen FPGAs.

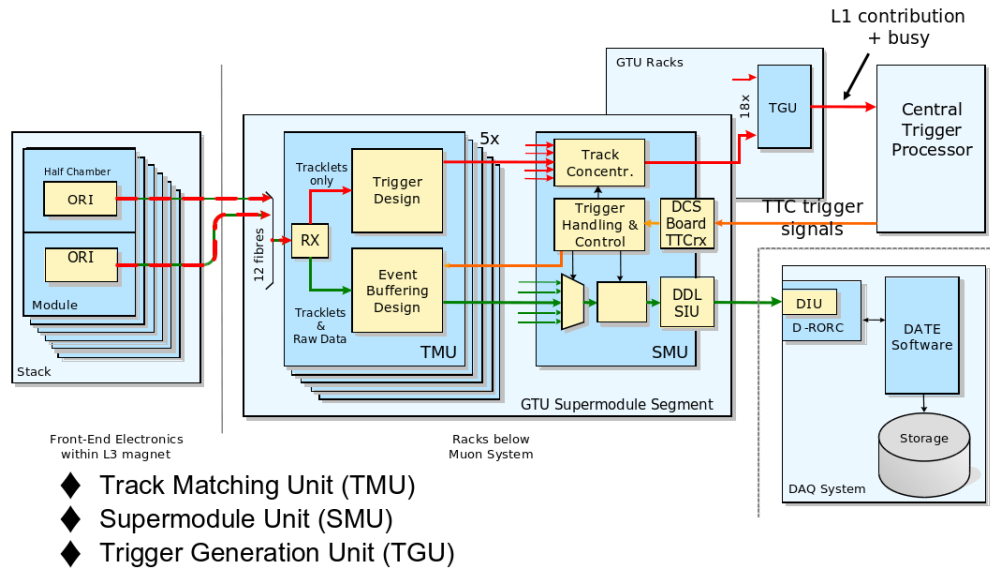


Abb. 4.9.: Der schematische Aufbau der GTU. Jeder Stapel im TRD ist über zwölf optische Fasern mit einer TMU verbunden. In der TMU werden die Spursegmente zusammengefügt und an je eine SMU pro Supermodul weitergegeben. Die SMU gibt die Spur-Daten an die TGU weiter, die eine Entscheidung über den L1-Beitrag trifft und an den CTP weitergibt. Die SMU verarbeitet außerdem die eingehenden Trigger des CTP, kontrolliert die TMUs und schickt die zwischengespeicherten Rohdaten bei einem L2A an die DAQ. Bild aus [Ang08]

Die GTU ist über zwei optische Verbindungen mit einer Transferrate von $2,6 \text{ GB s}^{-1}$ pro Modul mit dem TRD verbunden. Je eine TMU empfängt die Spursegmente eines Stapels, rekonstruiert die Spuren und berechnet die Impulse der Teilchen. Diese Informationen werden über die 18 SMUs an die TGU weitergegeben, der eine L1-Beitrag-Entscheidung trifft und $6,1 \mu\text{s}$ nach der Kollision an den CTP weitergibt. Dazu werden Teilchen mit hohen Impulsen ausgesucht. Über Koinzidenzen zwischen zwei Supermodulen können außerdem Di-Leptonen identifiziert werden. Des Weiteren gibt die TGU das busy-Signal weiter.

Wird von dem CTP ein L1A ausgegeben, werden die Rohdaten der MCMs über die optischen Verbindungen ausgelesen und in den SMUs zwischengespeichert. Dort bleiben sie bis der CTP einen L2A, der den Abtransport zum DAQ initiiert, oder einem L2R, der die Daten verwerfen lässt, schickt. Wenn $500 \mu\text{s}$ nach der Kollision weder L2A noch L2R gesendet wurde, werden die Daten ebenfalls verworfen. Um die Totzeit des TRD zu verringern, können in den SMUs mehrere Ereignisse zwischengespeichert werden. Dadurch können neue Triggersequenzen gestartet werden, sobald der

Transport der Rohdaten vom Supermodul zu den SMUs abgeschlossen ist.

Die SMUs sind über ein DCS-Board mit dem DCS-Netzwerk verbunden. Darüber kann die GTU kontrolliert und konfiguriert werden.

In diesem Kapitel wurde der TRD von ALICE vorgestellt. Er ist segmentiert in 18 Supermodule, die in Münster montiert werden. Jedes Supermodul besteht aus 30 einzelnen Modulen, die die Spuren und gegebenenfalls die Übergangsstrahlung von Elektronen nachweisen.

Um eine schnelle Triggerentscheidung zu treffen, ist ein Großteil der Elektronik auf den Modulen montiert. Durch diesen massiven Einsatz paralleler Elektronik können mit dem TRD schon $6,5\ \mu\text{s}$ nach einer Kollision physikalisch interessante Ereignisse markiert werden.

Da die Elektronik die meiste Zeit in einem Ruhemodus ist, braucht der TRD den schnellen Pretrigger, der sie aufweckt. Insgesamt erwartet die Elektronik der Supermodule drei Triggerpulse in fest definierten Zeitfenstern.

5. Kosmische Strahlung

Bevor die Supermodule in ALICE eingebaut werden können, müssen sie in Münster getestet und kalibriert werden. Dazu wird eine Quelle hochenergetischer geladener Teilchen benötigt, deren Ionisationsspuren im Supermodul aufgenommen werden können. In Münster wird dazu natürliche kosmische Strahlung benutzt.

Kosmische Strahlung wird in primäre und sekundäre Strahlung unterteilt. Primäre kosmische Strahlung entstammt astrophysikalischen Quellen. Sie wechselwirkt mit der Erdatmosphäre und produziert dabei neue Teilchen. Diese werden als sekundäre kosmische Strahlung bezeichnet und durchdringen die Atmosphäre teilweise bis auf Meeresniveau. Die Spuren der geladenen Teilchen, die das Supermodul in Münster durchqueren, können aufgezeichnet werden und zur Kalibrierung benutzt werden.

5.1. Primäre kosmische Strahlung

Primäre kosmische Strahlung setzt sich aus Teilchen zusammen, die aus astrophysikalischen Quellen stammen und durch diese beschleunigt werden. Solche Quellen sind unter anderem Sonneneruptionen und Supernovae. Die genaue Herkunft ist jedoch insbesondere für hochenergetische Teilchen nicht vollständig geklärt. Die Lokalisierung der Quellen wird unter anderem dadurch erschwert, dass geladene Teilchen durch kosmische Magnetfelder abgelenkt werden.

Die Strahlung besteht hauptsächlich aus freien Protonen (79 %) und Heliumkernen (15 %) [Ams08]. Zusätzlich sind noch Elektronen und solche Atomkerne, die als Zwischenprodukte aus der Nukleosynthese hervorgehen, in geringem Umfang enthalten. Das Energiespektrum der primären kosmischen Strahlung erstreckt sich bis über 10^{20} eV, wobei die Intensität mit zunehmender Energie abnimmt. Sie kann bei geringeren Energien ($\lesssim 10^{14}$ eV) durch direkte Messungen, z.B. mit Wetterballons, nachgewiesen werden. Für größere Energien werden indirekte Messungen benutzt, die z.B. Luftschauer nachweisen, die von hochenergetischen kosmischen Teilchen erzeugt werden.[Gru00]

5.2. Sekundäre kosmische Strahlung

Primäre kosmische Strahlung wechselwirkt beim Eindringen in die Erdatmosphäre mit den Atomkernen. Die Strahlungslänge für Photonen und Elektronen in Luft ist $X_0 = 36,66 \text{ g cm}^{-2}$, während

die Wechselwirkungslänge für Hadronen $\lambda = 90.0 \text{ g cm}^{-2}$ beträgt. Bei einer Massenbelegung¹ der Atmosphäre von ca. 1000 g cm^{-2} entspricht das 27 Strahlungslängen für Photonen und Elektronen und 11 Wechselwirkungslängen für Hadronen. [Gru00]

Je nachdem welche Teilchen betrachtet werden, werden hadronische und/oder elektromagnetische Schauer erzeugt. Protonen, die häufigsten Teilchen der primären Strahlung, stoßen bei ca. 15 – 20 km Höhe mit den Atomkernen der Luft und erzeugen hauptsächlich Pionen, die leichtesten Mesonen, und Kaonen. Es werden ca. zehnmal so viele Pionen wie Kaonen produziert [Ams08]. Die geladenen Pionen zerfallen mit einer Wahrscheinlichkeit von mehr als 99% [PDG10] über den leptonenischen Zerfall in Myonen² und Neutrinos:

$$\pi^+ \rightarrow \mu^+ + \nu_\mu \quad (5.1)$$

$$\pi^- \rightarrow \mu^- + \bar{\nu}_\mu \quad (5.2)$$

K^+ und K^- zerfallen zu 63% ebenfalls in Myonen und Neutrinos oder mit einer Wahrscheinlichkeit von 28% über den hadronischen Zerfall in weitere Pionen [PDG10]. K_s^0 zerfallen in zwei Pionen, während K_l^0 in drei Pionen oder semileptonisch in ein Pion und ein Elektron oder Myon mit entsprechendem Neutrino zerfallen.

Die π^0 zerfallen zu mehr als 98% [PDG10] in zwei Photonen und lösen damit elektromagnetische Subschaue aus: durch Paarbildung werden aus den Photonen Elektron–Positron–Paare erzeugt, die wiederum Bremsstrahlung emittieren. Durch die geringere Strahlungslänge, werden diese Schauer im Vergleich zu hadronischen relativ schnell gestoppt. [Ams08]

Reicht die Energie des ursprünglichen Teilchens aus, gelangt ein Teil der sekundären Strahlung bis auf Meeresniveau.

Zusammensetzung der kosmischen Strahlung in der Atmosphäre und auf Meeresniveau

In Abb. 5.1 ist der vertikale Teilchenfluss gegen die Höhe für die häufigsten Teilchen aufgetragen. Dabei sind nur Teilchen mit Energien $E > 1 \text{ GeV}$ berücksichtigt.

Bei großen Höhen dominieren Protonen und Neutronen. Sie wechselwirken mit den Atomkernen der Erdatmosphäre und produzieren hadronische und elektromagnetische Schauer. Die Intensität vertikaler Protonen auf Meeresniveau mit $E > 1 \text{ GeV}$ ist ca. $0,9 \text{ m}^{-2}\text{s}^{-1}\text{sr}^{-1}$ [Ams08].

Die Elektronen und Photonen aus den elektromagnetischen Schauern werden aufgrund ihrer vergleichsweise geringen Strahlungslänge relativ schnell gestoppt und der Fluss hat ein Maximum bei einer Höhe von 15 km. Die Intensität vertikaler Elektronen auf Meeresniveau mit $E > 1 \text{ GeV}$ ist ca. $0,2 \text{ m}^{-2}\text{s}^{-1}\text{sr}^{-1}$. Zusätzlich gibt es noch Elektronen geringer Energie aus Myon–Zerfällen.

Der Myonen–Fluss fällt mit sinkender Höhe nur leicht ab. Myonen stammen hauptsächlich aus

¹ Die Massenbelegung entspricht der Masse, die in einer Fläche integriert über eine Strecke enthalten ist. In diesem Fall wird also die Masse über die Länge der Atmosphäre integriert.

² Myonen schließt hier und im Folgenden Antimyonen ein. Analog werden Elektronen und Positronen mit Elektronen abgekürzt.

Zerfällen von geladenen Mesonen in einer Höhe von ca. 15 km. Beim Durchqueren der Atmosphäre verlieren sie bis zum Meeresniveau über Ionisation ca. 2 GeV [Ams08]. Abhängig von ihrer Geschwindigkeit und der daraus resultierenden Zeitdilatation können sie auch in Elektronen und Neutrinos zerfallen. Ein 2,4 GeV Myon hat beispielsweise eine Zerfallslänge von 15 km, die durch den Energieverlust auf ca. 8,7 km reduziert wird. Die Winkelverteilung von Myonen auf Meeresniveau ist $\propto \cos^2 \theta$. [Ams08].

80 % aller geladenen kosmischen Teilchen auf Meeresniveau sind Myonen. Sie haben eine durchschnittliche Energie von ca. 4 GeV [Gru00]. Die Intensität vertikaler Myonen auf Meeresniveau mit $E > 1$ GeV ist ca. $70 \text{ m}^{-2}\text{s}^{-1}\text{sr}^{-1}$ [Ams08].

Neutrinos sind auf Meeresniveau noch zahlreicher als Myonen vorhanden. Sie können die Atmosphäre nahezu ungehindert durchqueren, da ihr Wirkungsquerschnitt klein ist.

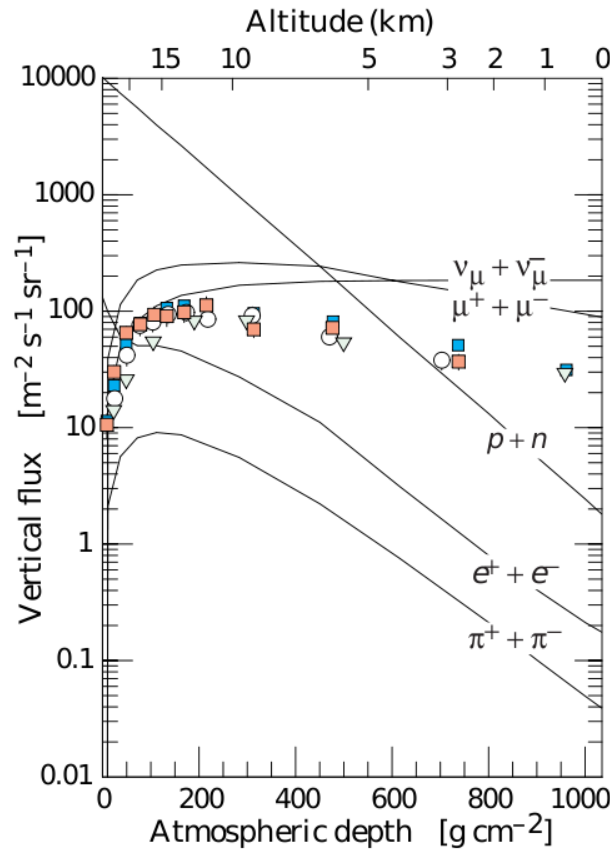


Abb. 5.1.: Der vertikale Fluss (logarithmisch) der häufigsten Teilchen gegen die Höhe in der Atmosphäre (logarithmisch) bzw. die Massenbelegung aufgetragen. Protonen und Neutronen dominieren in großen Höhen, werden aber durch Stöße mit Kernen der Luftatome gestoppt und produzieren sekundäre kosmische Strahlung. Auf Meeresniveau dominieren Myonen und Neutrinos. Die Punkte zeigen Messungen negativer Myonen. Bild aus [Ams08]

Für den Test der Supermodule in Münster (ca. 60 m über NN) ist der Teil der sekundären Teilchen, die eine elektrische Ladung tragen, geeignet. Es können also hauptsächlich die Spuren von

Myonen, aber auch in geringerem Umfang die Spuren von Protonen und Elektronen zum Testen und Kalibrieren der Supermodule aufgenommen werden.

6. Detektion kosmischer Teilchen in Münster

In diesem Kapitel wird zunächst erläutert, warum Spuren von kosmischen Teilchen aufgenommen werden und wie sie zur Kalibrierung des Supermoduls benutzt werden können.

Um kosmische Teilchen mit den Supermodulen aufnehmen zu können, müssen sie zunächst mit anderen Detektoren nachgewiesen werden, damit ein Trigger für das Supermodul erzeugt werden kann. Dazu wurde in [Bat07] ein Koinzidenztrigger entwickelt, der aus zwei Szintillatorebenen über und unter dem Supermodul besteht. Dieser Aufbau wurde teilweise übernommen und im Rahmen dieser Arbeit um das *V1495 General Purpose Board* von Caen ergänzt.

Im Anschluss an den Aufbau wird eine Messung vorgestellt, bei der das V1495 benutzt wird, um die Ausleseelektronik der Szintillatoren zu kalibrieren. Außerdem wird die Koinzidenzbedingung zwischen den beiden Szintillatorebenen überprüft und eingestellt.

6.1. Motivation für einen Trigger auf kosmische Teilchen

Wie in Kapitel 5 diskutiert, besteht der Großteil der kosmischen Strahlung auf Meeresniveau aus Myonen und Anti-Myonen mit einer durchschnittlichen Energie von 4 GeV. Myonen und Antimyonen sind Leptonen mit der Ladung $\pm e$ und haben eine Masse von ca. $106 \text{ MeV } c^{-2}$. Wenn sie die TRD-Module durchqueren, ionisieren sie das Gas und ihre Spuren können ausgelesen werden.

Myonen erzeugen keine Übergangsstrahlung, da sie bei 4 GeV einen Lorentzfaktor $\gamma \approx 38$ haben. Außerdem ist die Orientierung der Module für kosmische Teilchen falsch, da diese erst das Driftvolumen und dann den zugehörigen Radiator durchqueren.

Für die Kalibrierung sind die Ionisationsspuren jedoch gut geeignet. Im Folgenden werden einige Möglichkeiten vorgestellt, wie die Spuren zur Kalibrierung des Supermoduls genutzt werden können.

Alignment – Ausrichtung der Module

Wenn die Module in die Supermodule eingebaut werden, lassen sie sich um wenige Millimeter verschieben, damit Ungenauigkeiten bei den Bohrungen oder der Supermodulhülle ausgeglichen werden können. Damit die Module möglichst mittig liegen, werden sie per Augenmaß justiert.

Die angestrebte Ortsauflösung des TRD liegt im Bereich von einigen Hundert Mikrometern. Damit diese erreicht werden kann, müssen die relativen Verschiebungen der Module zueinander bekannt sein.

Die Spuren der kosmischen Teilchen sind dazu hervorragend geeignet, da sie nicht durch ein Magnetfeld abgelenkt werden und somit gerade sind. Um die Verschiebungen zu bestimmen, wird die Position von einem Modul eines Stapels als fix betrachtet und die Verschiebung der anderen Module relativ zu diesem mit den geraden Spuren bestimmt.

Werden die Spuren mit ausreichender Statistik ausgewertet, können die Verschiebungen exakt berechnet und bei der Rekonstruktion der Daten von ALICE berücksichtigt werden. Für weitere Informationen sei auf [Sic09] und [Gat10] verwiesen.

Kalibrierung der Verstärkung durch die PASA

Der Vorverstärker (PASA) in den Multi-Chip-Modulen liest die Ladung aus, die auf den Kathodenpads induziert wird, und verstärkt sie (vgl. Kap. 4.6). Diese Verstärkung ist nicht gleich für die unabhängigen PASA. Deshalb muss für jeden Kanal ein individueller Korrekturfaktor eingeführt werden, der die Verstärkung angleicht. Dieser Vorgang wird *Gain Calibration* genannt.

Wenn kosmische Teilchen mit einer ausreichenden Statistik aufgenommen werden, also jedes Pad oft getroffen wurde, ist die integrierte deponierte Ladung pro Kathodenpad gleich. Indem die gemessenen, deponierten Ladungen verglichen werden, können die Korrekturfaktoren für jeden Kanal bestimmt werden.

Kritisch für diese Kalibrierung ist die erreichte Statistik. Mit dem Trigger müssen also möglichst viele Teilchen identifiziert werden. In [Alb10] ist eine detaillierte Diskussion der Gain Calibration zu finden.

Weitere Möglichkeiten

Zusätzlich zur Gain Calibration und dem Alignment werden durch die Aufzeichnung von kosmischen Teilchen weitere Möglichkeiten eröffnet:

- Die komplette Auslekette von der Ionisationsspur zur DAQ wird mit echten Daten getestet. Dadurch wird das Supermodul unter nahezu identischen Bedingungen wie am LHC getestet und es kann garantiert werden, dass die Detektorelektronik korrekt arbeitet.
- Die Verarbeitung von echten Daten ermöglicht es, andere Elemente des TRD zu testen. So können z.B. neue Versionen der GTU getestet werden.

6.2. Aufbau zur Detektion kosmischer Teilchen

In [Bat07] wurde ein Koinzidenztrigger entwickelt, mit dem kosmische Teilchen detektiert werden können. Er besteht aus zwei Szintillator-Ebenen, die unter und über dem Supermodul platziert sind, und einer Koinzidenzeinheit, die die Signale der beiden Ebenen auf Koinzidenz prüft. Für das neue Triggersystem wurden die Szintillatoren und Teile der Ausleseelektronik übernommen und um ein *V1495 General Purpose VME Board* von Caen ergänzt.

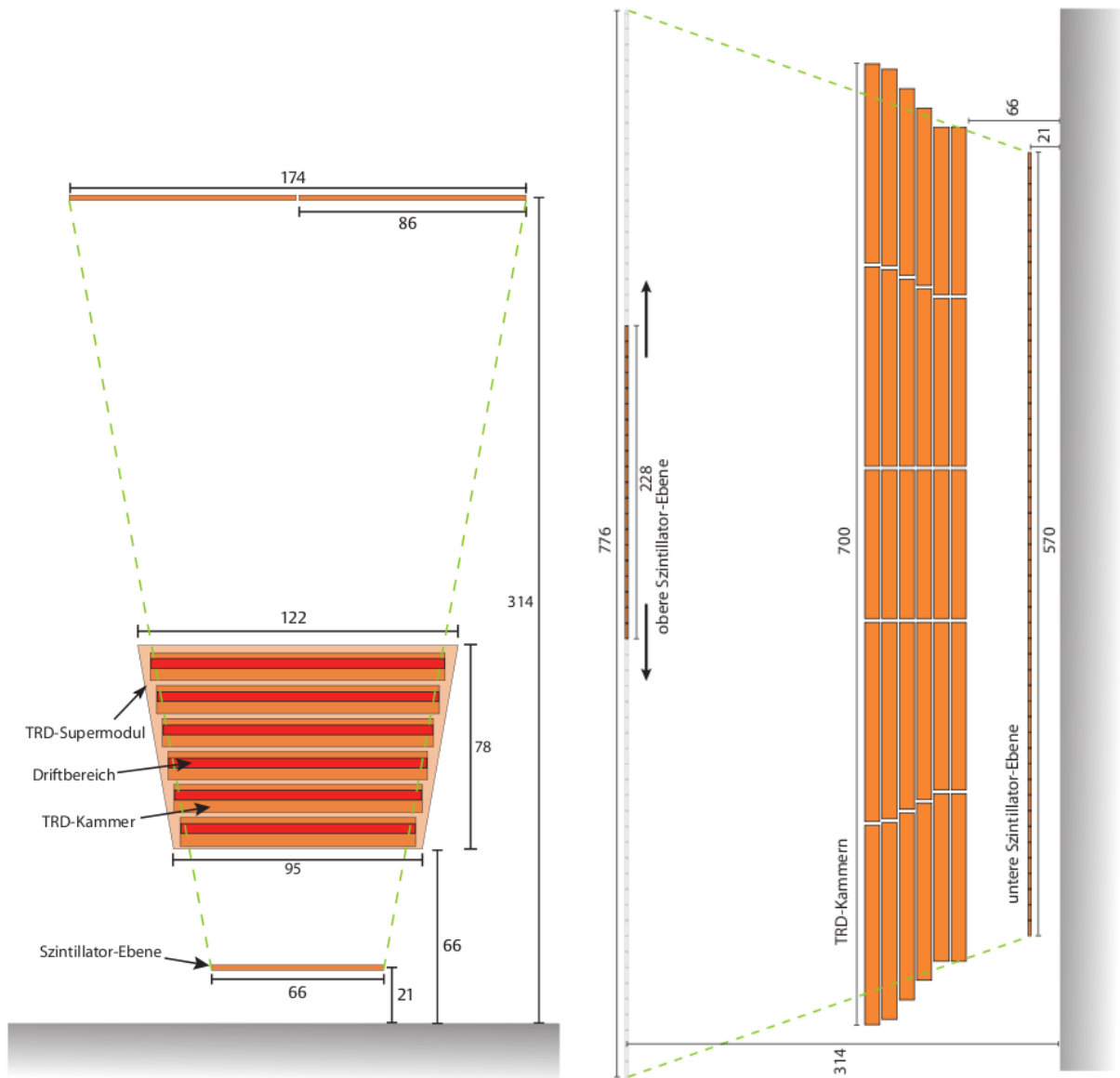


Abb. 6.1.: Schematischer Aufbau des Teststandes von vorne (links) und von der Seite (rechts). Die Szintillatoren sind so ausgelegt, dass ein Teilchen, das in beiden Ebenen ein Signal erzeugt (Koinzidenz), auch das Supermodul durchquert. Alle Maße sind in cm. Bilder aus [Bat07]

In Abb. 6.1 ist der schematische Aufbau der Szintillatorebenen mit dem Supermodul von vorne und von der Seite gezeigt. Das Supermodul liegt auf einem Metallgestell ca. 66 cm über dem Boden. Über und unter dem Supermodul ist eine Szintillatorebene mit 40 bzw. 50 unabhängigen Szintillatoren installiert, die jeweils von einem Photovervielfacher ausgelesen werden.

Die obere Ebene (Top) hat insgesamt eine aktive Fläche von $3,8 \text{ m}^2$ und kann über eine Schiene auf einer Länge von 7,76 m bewegt werden, während die untere Ebene (Bottom) eine aktive Fläche von $3,6 \text{ m}^2$ abdeckt und fest installiert ist. Die aktiven Flächen und die maximale Auslenkung

der oberen Ebene wurde so gewählt, dass ein kosmisches Teilchen, das beide Ebenen durchquert, auch durch das Supermodul fliegen muss. Das ist in den Zeichnungen durch die gestrichelten Linien angedeutet.

Aufgrund der Trapezform des Supermoduls, muss die obere Ebene breiter als die untere sein, damit diese Bedingung erfüllt ist. Um die komplette obere Ebene abzudecken, würden ca. 3,5 mal so viele Szintillatoren benötigt. Die kleinere bewegliche Ebene stellt einen Kompromiss zwischen erreichbarer Statistik pro Zeit und Kosten bzw. Aufwand dar. Indem die obere Ebene beweglich ist, kann trotzdem die ganze Akzeptanz abgedeckt werden.

Jeweils zehn Szintillatoren werden zu einer Gruppe zusammengefasst und mit T1–T4 für die obere Ebene und B1–B5 für die untere Ebene bezeichnet. Jede Gruppe in Bottom deckt ungefähr einen Stapel des Supermoduls ab.

In Abb. 6.2 ist die Verarbeitung der Szintillatorsignale schematisch dargestellt. Jeder Szintillator wird über einen Photovervielfacher ausgelesen. Die Signale einer Gruppe werden in einem Trefi-Modul gesammelt. Dort werden sie verstärkt, diskriminiert und mit einem logischen ODER verschaltet und dann als ein NIM-Signal¹ ausgegeben. Die Signale der neun Trefis werden wiederum im V1495 gesammelt und weiterverarbeitet².

Für die folgenden Testmessungen wurde ein Design für den FPGA des V1495 in VHDL entwickelt, mit dem es möglich ist, die Koinzidenzbedingung und die Diskriminatorschwellen der Trefis zu optimieren. Die einzelnen Bauteile und ihre Funktion werden im Folgenden beschrieben.

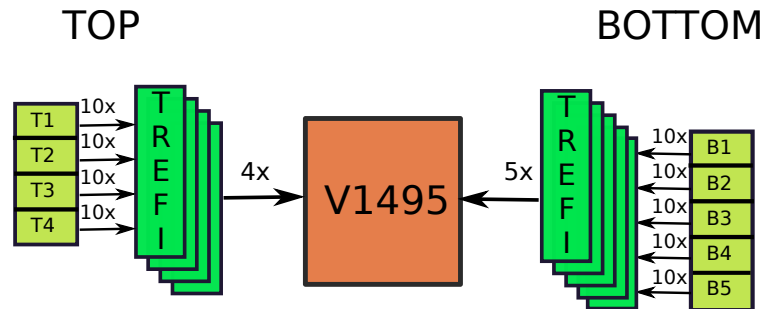


Abb. 6.2.: Signalverarbeitung der Szintillatoren. Die Signale von jeweils zehn Photovervielfachern werden in einem Trefi verstärkt, diskriminiert und zusammengefasst. Die Signale der neun Szintillatorgruppen werden im V1495 gesammelt verarbeitet.

6.2.1. Szintillatoren

Das Funktionsprinzip von Szintillatoren wurde bereits in Kap. 2.3 diskutiert. Bei den verwendeten Szintillatoren handelt es sich um Szintillatoren, die ursprünglich für das *Collider Detector at Fermilab*-Experiment am Tevatron hergestellt und für Münster auf die passenden Größen zugeschnitten

¹ In Anhang B ist eine Tabelle mit allen Logik-Standards, die im Aufbau benutzt werden, zu finden.

² Im alten Aufbau aus [Bat07] wurden die Signale der Trefis in einer Koinzidenzeinheit verarbeitet, um einen Pretrigger für das Supermodul zu generieren.

wurden.

Der Grundstoff der Szintillatoren ist Polystyrol, dem als primärer Fluoreszenzstoff 2 % Para-Terphenyl (p-TP) und als Wellenlängenschieber 0,03 % Diphenyloxazolybenzol (POPOP) beigemischt sind. Durchqueren hochenergetische geladene Teilchen den Szintillator, wird das Para-Terphenyl in angeregte Zustände versetzt. Diese zerfallen unter Emission von Photonen mit einer maximalen Wellenlänge $\lambda = 440 \text{ nm}$ und einer Abklingzeit von $\tau = 5 \text{ ns}$. Vom Wellenlängenschieber POPOP werden hauptsächlich Photonen mit $\lambda = 420 \text{ nm}$ emittiert. [Bat07]

Die Szintillatoren sind in Papier eingewickelt, um sie vor Licht zu schützen. So wird sichergestellt, dass nur Signale von kosmischen Teilchen und nicht von Photonen erzeugt werden. Außerdem wird damit die Reflektivität erhöht, so dass erzeugte Photonen den Szintillator nicht verlassen. Das eine Ende von jedem Szintillator wird mit 5-lagigem Teflon-Band abgedeckt, um eine hohe Reflektivität zu gewährleisten. Am anderen Ende wird ein Photovervielfacher angebracht, der die erzeugten Photonen nachweist und ein elektrisches Signal ausgibt.

Die Szintillatoren sind in neun Gruppen mit jeweils 10 Szintillatoren eingeteilt. In Anhang C ist die genaue Anordnung und Nummerierung der einzelnen Szintillatoren dargestellt, wie sie im Folgenden verwendet wird. Die Gruppen werden mit T1–T4 bzw. B1–B5 und einzelne Szintillatoren z.B. mit B3–8, also der achte Szintillator in Gruppe B3, beschrieben.

6.2.2. Photovervielfacher

Um die Signale der Szintillatoren auszulesen, werden FEU-183-1 Photovervielfacher des *Central Scientific Research Institute Electron* in St. Petersburg, Russland verwendet.

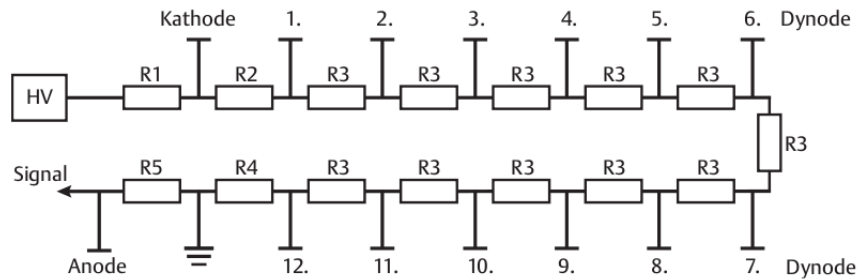


Abb. 6.3.: Das Schaltbild der Photovervielfacher. Die Widerstände sind: $R1 = 11 \text{ k}\Omega$, $R2 = 2,4 \text{ M}\Omega$, $R3 = 1,2 \text{ M}\Omega$, $R4 = 620 \text{ k}\Omega$, $R5 = 1,3 \text{ k}\Omega$ [Bat07]

In Abb. 6.3 ist das Schaltbild eines Photovervielfachers gezeigt. Jeweils zehn Vervielfacher werden über ein CAEN N470 Netzgerät mit einer negativen Hochspannung von ca. $1400 \text{ V} - 1500 \text{ V}$ versorgt. Indem die Dynoden mit Widerständen in Reihe geschaltet sind, wird ein Spannungspotential zwischen ihnen aufgebaut. Die Photokathode hat einen Durchmesser von 76 mm . Treffen Photonen darauf, lösen sie Elektronen aus, die über das Spannungspotential der 12 Dynoden beschleunigt und vervielfacht werden. An der Anode wird das Signal ausgelesen. [Bat07]

6.2.3. Trefi-08-Modul

Die Trefi-08-Module verstärken und diskriminieren die Signale von jeweils zehn Photoverstärkern und geben ein logisches Signal (NIM) aus.

Jedes Modul hat 16 Eingänge und acht Verstärker- und Diskriminator-Kanäle. Jeweils zwei Eingänge werden linear addiert, um den Faktor 10 verstärkt und dann diskriminiert. Durch die Diskrimination wird idealerweise das Rauschen abgeschnitten, aber keine echten Ereignisse unterdrückt. Die Schwelle für die Diskriminatoren kann über ein Potentiometer zwischen 20 mV und 4,2 V eingestellt werden. Sie kann nicht für jeden Kanal individuell eingestellt werden und ist somit für alle Kanäle innerhalb eines Moduls gleich.

Im normalen Betrieb werden die acht Kanäle mit einem logischen ODER verschaltet und über einen LEMO-Stecker als NIM-Signal ausgegeben. Dadurch reduziert sich die Zahl der zu verarbeitenden Szintillatorsignale von 90 auf neun, wobei immer noch eine ausreichende Segmentierung gegeben ist.

Zusätzlich wird jeder einzelne der acht Kanäle als ECL-Signal ausgegeben. Diese Ausgänge werden in der folgenden Messung genutzt, um die optimalen Diskriminatorschwellen unter Berücksichtigung aller Szintillatoren einzustellen. In Anhang C ist die Frontblende eines Trefi-Moduls mit allen Ein- und Ausgängen abgebildet.

6.2.4. CAEN V1495 General Purpose VME Board

Das *CAEN V1495 General Purpose VME Board* ist eine 1U breite VME-Karte, auf der zwei FPGAs installiert sind, mit denen diverse Ein- und Ausgänge angesprochen werden können. Es ist die zentrale Triggereinheit des neuen Triggersystems, das in Kap. 7 vorgestellt wird. Hier wird es zunächst benutzt, um die Diskriminatorschwellen der Trefis und die Koinzidenzbedingung einzustellen. In Abb. 6.4 ist ein Blockdiagramm des V1495 abgebildet.

FPGA Bridge

Der *FPGA Bridge* ist mit der VME Schnittstelle und über einen proprietären lokalen Bus mit dem zweiten *FPGA User* verbunden. Seine Aufgabe besteht darin, den *FPGA User* mit dem vom Benutzer bereitgestellten Design zu programmieren und ihn mit dem VME-Bus zu verbinden. Das aktuelle Design wird in einem Flash-Speicher abgelegt, wo es dauerhaft, also auch wenn das V1495 ausgeschaltet wird, gespeichert wird. Der *FPGA User* wird nach dem Anschalten automatisch mit dem Design aus dem Flash-Speicher programmiert.

FPGA User

Der *FPGA User* ist ein Altera Cyclone EP1C20, der vom Benutzer frei programmiert werden kann. Er hat 20060 Logikelemente³ und 294.912 RAM Bits. Zusätzlich sind zwei Phasenregelschaltungen

³ Ein Logikelement ist die kleinste logische Einheit, die auf einem FPGA implementiert ist.

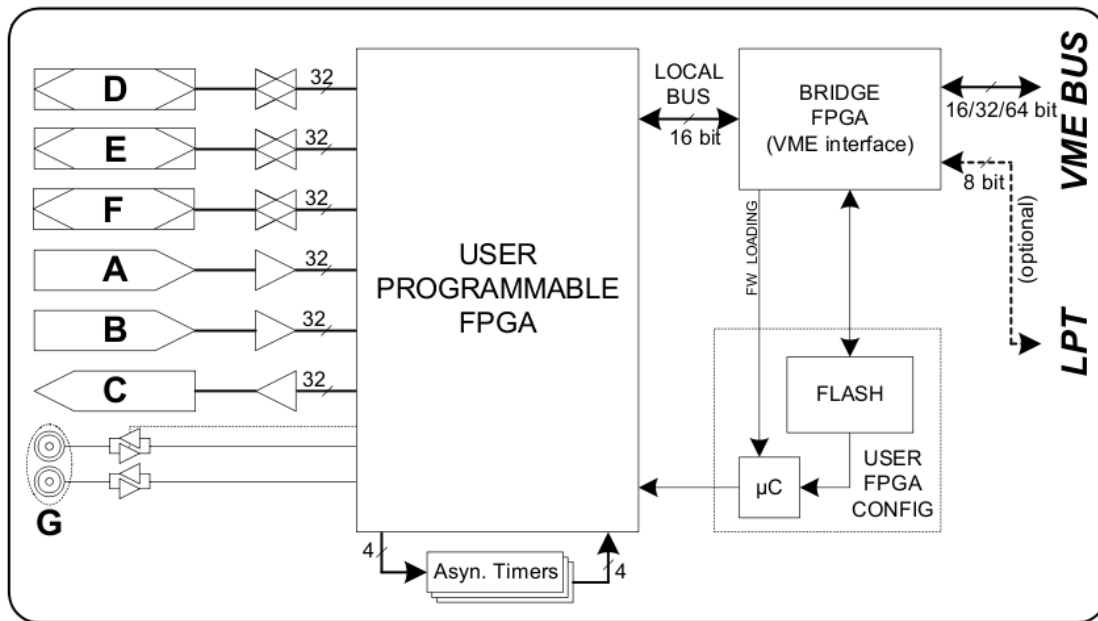


Abb. 6.4.: Ein Blockdiagramm des V1495. Das V1495 verfügt über einen FPGA, der frei programmiert werden kann und über einen zweiten FPGA auf die VME-Schnittstelle zugreifen kann. Außerdem steht eine Vielzahl an Ein- und Ausgängen zur Verfügung. Daneben sind noch vier asynchrone Timer vorhanden, die mit dem FPGA User kontrolliert werden. Bild aus [CAE10]

(PLL kurz für *Phase Locked Loop*) vorhanden, mit denen Taktgeber multipliziert, dividiert oder in der Phase verschoben werden können.

Taktgeber werden über ein globales Netzwerk mit acht separaten Kanälen an die Logikelemente, den RAM und die Ein- und Ausgangspins verteilt. Im FPGA steht ein interner 40 MHz Taktgeber zur Verfügung. Zusätzlich kann ein Taktgeber von außen eingespeist werden. [Alt03]

Der FPGA User kann außerdem auf vier asynchrone Timer (zwei *Programmable Delay Lines* (PDL) und zwei *Delay Line Oscillators* (DLO)) zurückgreifen, die zur Signalformation oder zur Verzögerung von Signalen genutzt werden können.

Die Designs für den FPGA User werden mit der Software *Quartus II Web Edition* von Altera in VHDL entwickelt und kompiliert. Bei allen Designs, die entwickelt wurden, wird die TTC-CLK (40,079 MHz) als Taktgeber benutzt. Alle Signale werden damit einsynchronisiert, außer es wird an anderer Stelle ausdrücklich erwähnt. Zur Synchronisation wird nach steigenden Flanken gesucht, indem das um einen Takt verzögerte Signal mit dem aktuellen Signal verglichen wird.

Ein- und Ausgänge

Mit FPGA User werden die Ein- und Ausgänge des V1495 kontrolliert. Diese unterteilen sich in vier vorinstallierte Blöcke (A, B, C und G). A und B haben jeweils 32 Eingänge die wahlweise LVDS-,

ECL- oder PECL-Logiklevel⁴ akzeptieren. C hat 32 LVDS-Eingänge und G bietet 2 Ein- oder Ausgänge mit TTL- oder NIM-Logik. Der Pin G0 ist außerdem dazu geeignet, einen Taktgeber von außen einzugeben. Hier wird die TTC-CLK eingespeist.

Zusätzlich wurde beim verwendeten V1495 noch eine Mezzanine-Erweiterung A395D angebracht, die wahlweise acht Ein- oder Ausgänge mit TTL- oder NIM-Logik bietet.[CAE10]

NIM-LVDS- und LVDS-NIM-Wandler

Da die bereits vorhanden Geräte größtenteils NIM-Signale als Ein- und Ausgänge verwenden, wurden in der Elektronikwerkstatt des Instituts für Kernphysik Münster ein NIM-LVDS- und ein LVDS-NIM-Wandler gebaut. Sie werden an Block B bzw. Block C angeschlossen und konvertieren jeweils 32 Kanäle. Da der TTCex, der die Triggersignale in optische Signale wandelt und weitergibt, nur ECL-Signale akzeptiert, werden bei dem LVDS-NIM-Konverter vier Ausgänge nicht in NIM, sondern in ECL gewandelt. Bei dem NIM-LVDS-Wandler muss beachtet werden, dass die Signale negiert werden. Deshalb müssen sie im FPGA erneut negiert werden.

Im Folgenden werden die Konverter nicht immer explizit erwähnt, sondern wie direkte Eingänge des V1495 behandelt. Die Schaltpläne der Konverter befinden sich in Anhang D.

6.3. Koinzidenzbedingung

Um einen Trigger für das Supermodul zu generieren, wenn es von einem kosmischen Teilchen durchquert wird, können die Signale der beiden Szintillatorebenen auf Koinzidenz geprüft werden. Wird in beiden Ebenen gleichzeitig ein Signal erzeugt, dann ist die Wahrscheinlichkeit sehr hoch, dass ein kosmisches Teilchen durch beide Ebenen und somit auch durch das Supermodul geflogen ist.

Damit die Signale aber auf Gleichzeitigkeit geprüft werden können, müssen die Signallaufzeiten berücksichtigt und angeglichen werden. Die Verarbeitungszeiten der Elektronik sind für beide Ebenen gleich, da für beide die gleiche Elektronik verwendet wird. Die Kabellängen zwischen den Photovervielfachern und dem V1495 sind jedoch verschieden. Hinzu kommt, dass die kosmischen Teilchen ca. $11_{-1,3}^{+3,6}$ ns [Bat07] für die Strecke zwischen Top und Bottom benötigen. Um diese Differenzen auszugleichen, müssen die Signale entsprechend verzögert werden.

Aufbau zur Bestimmung der Verzögerung

In Abb. 6.5 ist der Aufbau skizziert, um die optimale Verzögerung zu bestimmen. Dazu wird eine Koinzidenzkurve aufgenommen, bei der die Koinzidenzrate kosmischer Teilchen zwischen Top und Bottom in Abhängigkeit der Verzögerung aufgezeichnet wird. Die Signale der oberen Ebene und die der unteren Ebene werden jeweils mit einem logischen ODER verknüpft. Sie werden zuvor nicht

⁴ In Anhang B ist eine Tabelle mit allen Logik-Standards, die im Aufbau benutzt werden, zu finden.

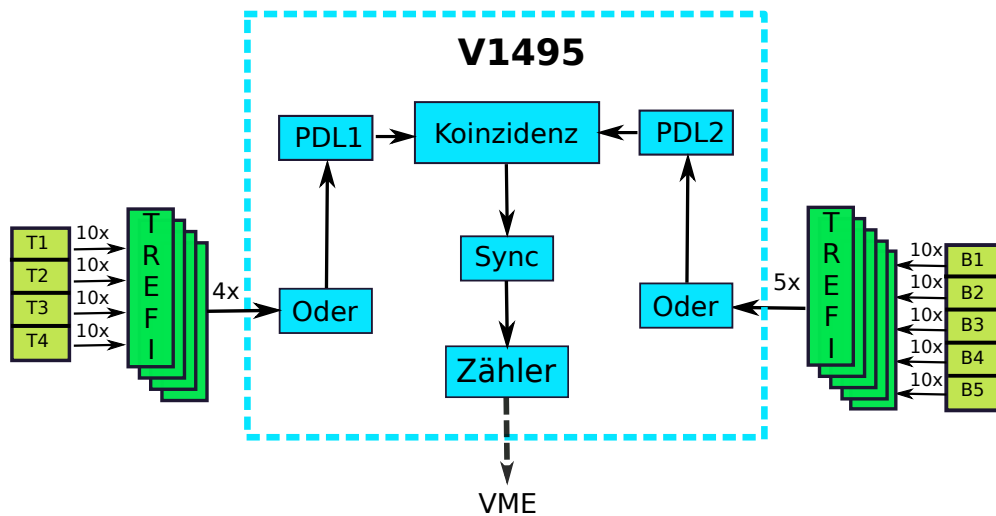


Abb. 6.5.: Um die Koinzidenzkurve aufzunehmen, werden alle Signale von TOP mit einem ODER jeweils zu einem Signal zusammengefasst und über eine PDL verzögert. Mit Bottom wird analog verfahren. Anschließend wird die Anzahl der Koinzidenzen der beiden Signale mit verschiedenen Verzögerungen gemessen und gezählt.

einsynchronisiert, da die Zeitinformation für die Koinzidenzkurve verloren gehen würde. Durch das Einsynchronisieren mit einem 40 MHz-Taktgeber wird ein Signal um bis zu 25 ns verschoben⁵.

Mit den beiden Programmable Delay Lines des V1495 können die Signale für Top und Bottom um $T_{\text{delay}} = T_{\text{offset}} + T_{\text{set}}$ verzögert werden. Dabei ist $T_{\text{offset}} = 32 \pm 2$ ns und T_{set} kann über ein Register in 1 ns-Schritten eingestellt werden.

Das logische UND der beiden verzögerten Signale wird mit der TTC-CLK einsynchronisiert, so dass es genau einen Takt lang hoch ist, wenn eine Koinzidenz gefunden wurde. Die Koinzidenzen werden mit einem Zähler für jede Verzögerung 60 Sekunden lang gezählt und dann über den VME-Bus ausgelesen.

Erwartung für die Messung

Im Idealfall sollte eine Rechteckfunktion mit einer Breite von 100 ns zu sehen sein, da die Signale der Trefis 50 ns lang sind und sich bei gleichzeitigen Ereignissen auf einer Länge von maximal 100 ns überschneiden. Die Rechteckfunktion wird jedoch durch verschiedene Effekte abgerundet und ausgeschmiert:

- Die Laufzeit des Lichtes im Szintillator beträgt bis zu 6 ns. [Bat07]
- Die Trefis diskriminieren das Signal der Photovervielfacher auf Grund ihrer Amplitude. Bei

⁵ Um asynchrone Signale in der fertigen Triggereinheit zu vermeiden, werden die Signale der Szintillatoren mit einem schnelleren Taktgeber einsynchronisiert und verglichen und erst dann mit der TTC-CLK synchronisiert (siehe 8.1.1). Diese Messung wurde jedoch durchgeführt als die Triggereinheit noch nicht fertig programmiert war.

schwachen Signalen ist die Steigung des Signals flacher und erreicht deshalb später die Diskriminatorschwelle. Dieser Effekt nennt sich Walk-Effekt und beträgt im Trefi ± 3 ns. [Bat07]

- Die Zeit, die ein kosmisches Teilchen braucht, um die Strecke von Top bis Bottom zurückzulegen, ist abhängig vom Einfallswinkel.

Ergebnisse der Messung

In Abb. 6.6 ist links die Koinzidenzkurve gezeigt. Die Verzögerung Δt ist so definiert:

$$\Delta t = \Delta t(\text{Top}) - \Delta t(\text{Bottom}) \quad (6.1)$$

Dabei ist $\Delta t(\text{Top})$ die Verzögerung, die für die obere Ebene und $\Delta t(\text{Bottom})$ die Verzögerung, die für die untere Ebene mit den PDLs eingestellt wird.

In dem Graph (Abb. 6.6 links) ist die Halbwertsbreite eingezeichnet, um den Mittelpunkt des Plateaus zu bestimmen. Dieser liegt bei ca. 9,3 ns. Um die Höhe des Plateaus zu bestimmen, wurde der Mittelwert aus den Werten im Plateau, also zwischen -5 ns und 25 ns, gebildet. Die Schnittpunkte der Halbwertsbreite mit der Kurve wurden bestimmt, indem die Kurve zwischen den Punkten linear angenähert wurde. Da $\Delta t = 0$ gut im Plateau liegt, ist es eigentlich nicht nötig, die Kabellängen anzupassen, um den Plateaumittelpunkt dorthin zu verschieben.

Im fertigen Design für die Triggereinheit wird zur Optimierung der Koinzidenzbedingung unter anderem die Pulslänge von Top und Bottom minimiert, um möglichst wenige zufällige Koinzidenzen zu erhalten (vgl. Kap. 8.1.1). Dadurch verringert sich die Breite des Plateaus. Deswegen wurden die Kabellängen so angepasst, dass der Plateaumittelpunkt bei $\Delta t \approx 0$ liegt.

Mit den neuen Längen wurde eine weitere Koinzidenzkurve aufgenommen, die in Abb. 6.6 rechts zu sehen ist. Hier wurde analog die Halbwertsbreite bestimmt und eingetragen. Der Mittelpunkt des Plateaus liegt mit den neuen Kabellängen bei $\Delta t = 1,2$ ns.

Die Kabel von Top entsprechen jetzt einer Verzögerung von 64 ns und die von Bottom einer Verzögerung von 52 ns. Das bestätigt gut die erwartete Flugzeit der kosmischen Teilchen zwischen Top und Bottom von $11^{+3,6}_{-1,3}$ ns.

6.4. Diskriminatorschwellen der Trefi-Module

In diesem Unterkapitel wird eine Messung vorgestellt, mit der die optimale Einstellung für die Diskriminatorschwellen der Trefis bestimmt wurde. Mit der Diskriminatorschwelle werden alle Signale unterdrückt, deren Amplitude unter der Schwelle liegen. Dadurch soll Rauschen abgeschnitten werden. Die Schwelle darf jedoch nicht zu hoch gewählt werden, da sonst auch die Signale von echten Ereignissen unterdrückt werden.

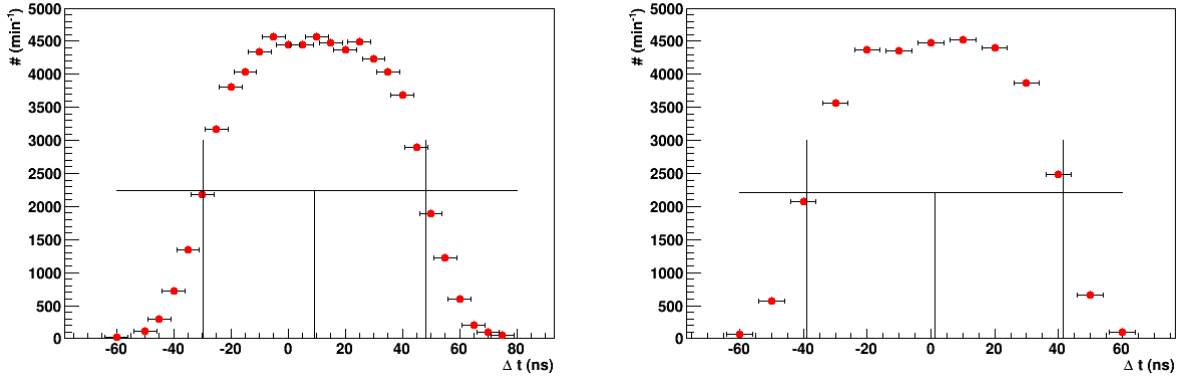


Abb. 6.6.: Koinzidenzkurven vor (links) und nach Anpassen (rechts) der Kabellängen. Die erwarteten Rechteckfunktionen werden durch verschiedene Effekte abgerundet und ausgeschmiert (siehe Text). Über die Halbwertsbreite wurde die Mitte des Plateaus bestimmt.

Die Schwellen der Trefis können nicht individuell für jeden Szintillator, sondern nur für jede 10er-Gruppe eingestellt werden. Die verwendeten Szintillatoren und Photovervielfacher weisen jedoch verschiedene Verhalten bei der Lichtausbeute bzw. Verstärkung auf. Bei der Installation wurde versucht ähnliche Szintillatoren und Photovervielfacher in einer Gruppe anzuordnen, um die Abweichungen innerhalb einer Gruppe zu minimieren [Bat07].

Um die Schwellen der Trefis einzustellen, wurde in [Bat07] jeweils der Szintillator einer Gruppe ausgewählt, der das schwächste Signal liefert und für verschiedene Diskriminatorschwellen vermessen. Auf dieser Grundlage wurde die Schwelle eingestellt. Mit dem V1495 ist es möglich, viele Signale gleichzeitig zu verarbeiten. Deshalb wurde eine neue Kalibrierung durchgeführt, bei der alle Szintillatoren berücksichtigt werden.

Dazu wurde ein Design für den FPGA User und ein Aufbau entwickelt, mit dem jeder Szintillator der oberen Ebene einzeln in Koinzidenz mit der gesamten unteren Ebene und analog jeder Szintillator der unteren Ebene einzeln mit der gesamten oberen Ebene vermessen werden kann. Als Teilchenquelle wird auch hier kosmische Strahlung benutzt.

Dadurch kann für jeden Szintillator einzeln ein Schwellenbereich bestimmt werden, in dem das Rauschen, aber keine echten Ereignisse abgeschnitten werden. Indem man alle Bereiche einer Gruppe berücksichtigt, erhält man einen Bereich, in dem alle Szintillatoren optimal arbeiten.

6.4.1. Aufbau für die Messung

Ziel ist es, die Koinzidenzrate zwischen einem Szintillator einer Ebene mit der kompletten anderen Ebene bei verschiedenen Diskriminator-Schwellen der Trefis aufzunehmen. Diese Messung soll für jeden Szintillator durchgeführt werden.

In Abb. 6.7 ist der schematische Aufbau gezeigt, mit dem die Koinzidenzrate des Szintillators T1-1 mit der unteren Szintillatorebene bei verschiedenen Diskriminatorschwellen aufgenommen werden

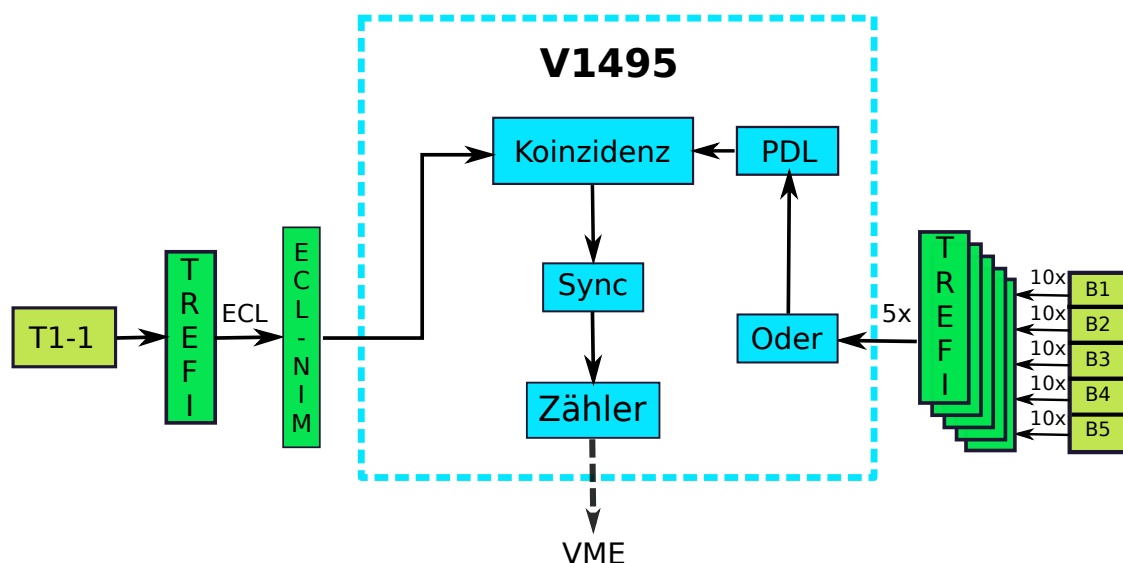


Abb. 6.7.: Hier ist der Aufbau zur Kalibrierung der Diskriminatorschwellen beispielhaft für den Szintillator T1-1 gezeigt. Die PDL wird hier benutzt, um die Verzögerung durch das lange ECL-Kabel auszugleichen. Um die Messzeit zu verkürzen, ist die Logik in zehnfacher Ausführung implementiert. So können zehn Szintillatoren auf einmal vermessen werden.

kann. Für alle anderen Szintillatoren ist der Aufbau analog.

Da diese Messung für jeden Szintillator gemacht werden soll, wird die Messzeit verkürzt, indem mehrere Szintillatoren gleichzeitig vermessen werden. Pro Trefi stehen acht unabhängige Verstärker- bzw. Diskriminatorkanäle zur Verfügung. Das bedeutet, dass maximal acht Szintillatoren pro Trefi gleichzeitig vermessen werden können. Deshalb wird jedes Trefi-Modul zweimal mit jeweils fünf angeschlossenen Szintillatoren vermessen. Die einzelnen Kanäle werden über den ECL-Ausgang der Trefis ausgegeben.

Mit einem Kabel, das in der Elektronikwerkstatt des IKP Münster gefertigt wurde, können die ECL-Signale von zwei Trefis gleichzeitig ausgelesen und über einen ECL-NIM-Konverter an das V1495 weitergegeben werden. So können die Kurven von zehn Szintillatoren parallel aufgenommen werden.

6.4.2. Signallaufzeiten für Koinzidenzbedingung

Damit das Signal eines Szintillators mit der kompletten anderen Ebene auf Koinzidenz geprüft werden kann, müssen die Signallaufzeiten, insbesondere des langen ECL-Kabels, berücksichtigt werden. Um die Differenz der Laufzeiten zu bestimmen, wurden zwei Koinzidenzkurven wie in 6.3 aufgenommen.

In Abb. 6.8 sind die aufgenommenen Koinzidenzkurven zu sehen. Links wird das Signal von acht Szintillatoren von Top über das ECL-Kabel mit Bottom auf Koinzidenz geprüft. Rechts wird das ECL-Kabel für acht Szintillatoren von Bottom benutzt.

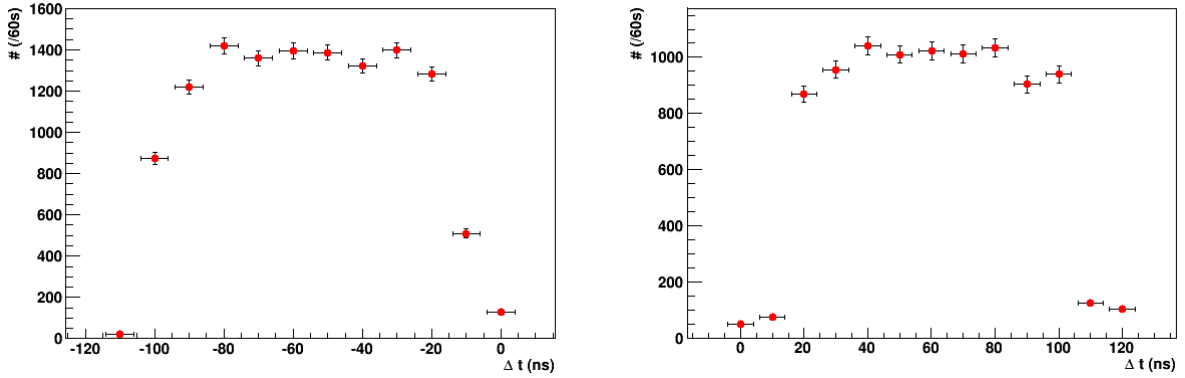


Abb. 6.8.: Hier ist die Anzahl der Koinzidenzen gegen die Verzögerung durch die PDLs aufgetragen. Δt ist die Differenz zwischen der Verzögerung der oberen Lage und der Verzögerung der unteren Ebene. Links: Koinzidenzkurve zwischen acht Szintillatoren der oberen Ebene (T1-7 – T1-10 und T4-7 – T4-10), die über das ECL-Kabel mit dem V1495 verbunden sind und der kompletten unteren Ebene (B1 – B5). Rechts: Koinzidenzkurve zwischen acht Szintillatoren der unteren Ebene (B3-7 – B3-10, B4-7 – B4-10), die über das ECL-Kabel mit dem V1495 verbunden sind und der kompletten oberen Ebene (T1 – T4).

Da es für die Messung lediglich wichtig ist, dass die Verzögerung im Plateau liegt, wird für die Bestimmung der Diskriminatorschwellen eine Verzögerung von 60 ns bzw. -60 ns eingestellt.

6.4.3. Messung

Um die Diskriminatorschwellen zu bestimmen, wurde die Koinzidenzrate kosmischer Teilchen von jedem Szintillator für 90 Sekunden bei verschiedenen Schwellen gemessen. Die Dauer der Messung ist ein Kompromiss zwischen erreichter Statistik und Messaufwand, da für jeden Szintillator ca. 15 bis 20 Zählraten bei verschiedenen Schwellen aufgenommen werden sollen und die Messung nicht automatisiert werden kann, weil die Diskriminatorschwellen mit einem Schraubendreher eingestellt werden müssen, während der aktuell eingestellte Wert mit einem Voltmeter beobachtet werden kann.

Um vergleichbare Werte zu erhalten wurde Top immer mittig über den zu vermessenden Szintillatoren von Bottom platziert. Um die oberen Szintillatoren zu vermessen, wurde Top mittig über Bottom platziert.

Erwartung

Für die aufgenommenen Kurven wird erwartet, dass die Rate bei niedrigen Schwellen am höchsten ist und dann abnimmt, bis das Rauschen abgeschnitten wird. Danach sollte ein Plateau zu beobachten sein, bis die Rate wieder sinkt, weil echte Signale unterdrückt werden.

Ergebnisse

In Abb. 6.9 sind beispielhaft die Koinzidenzen zwischen dem Szintillator B2-1 und Top für verschiedene Diskriminatorschwellen gezeigt. Hier ist die erwartete Kurve sehr gut zu beobachten. Wenn die Schwelle ausreicht, um das Rauschen abzuschneiden, bildet sich ein Plateau aus. Danach fällt die Rate wieder, weil echte Koinzidenzen abgeschnitten werden. Die vertikalen Striche markieren das Plateau. Die Grenzen wurden hier und für alle anderen Szintillatoren für das deutlich erkennbare Plateau gewählt. Wenn die Entscheidung nicht eindeutig ist, wurde immer der konservativere Wert gewählt.

In Abb. 6.10, 6.11 und 6.12 sind die Koinzidenzraten von jedem Szintillator gegen die verschiedenen Diskriminatorschwellen aufgetragen. Sie sind dabei immer in den entsprechenden 10er-Gruppen zusammengefasst.

Es lässt sich erkennen, dass die Kurven der Erwartung entsprechen. Bei fast allen Kurven bildet sich nach einem anfänglichen Abfall der Rate ein Plateau aus. Nach dem Plateau schließt sich der Bereich an, in dem die Raten sinken, weil echte Signale von kosmischen Teilchen abgeschnitten werden.

Teilweise sind Ausreißer zu beobachten, die nicht der Erwartung entsprechen. Es ist jedoch zu berücksichtigen, dass die Messung mit kosmischen Teilchen gemacht wurde und der Teilchenstrom nicht konstant ist. Aufgrund der relativ geringen Statistik können Schwankungen im Teilchenstrom bzw. dem Energiespektrum zu diesen Ausreißern führen.

Außerdem hat es sich als schwierig herausgestellt, die Schwellen mit den Potentiometern der Trefis einzustellen. Zum Beispiel wurden Sprünge der Schwelle bei minimalen Veränderungen der Stellschraube beobachtet. Für bestimmte Bereiche (abhängig vom Trefi) ließ sich des Weiteren kein konstanter Wert einstellen.

Obwohl die Form der Kurve für die meisten Szintillatoren der Erwartung entspricht, lässt sich auch beobachten, dass sie sehr verschiedene Werte für die Plateau-Grenzen aufweisen. Um die beste Diskriminatoreinstellung für eine Gruppe zu bestimmen, wurden die Grenzen für jeden Szintillator abgeschätzt. Im Anschluss wurde für jede Gruppe die maximale untere Grenze und die minimale obere Grenze für das Plateau ermittelt. Im Bereich zwischen diesen beiden Werten liegen alle Szintillatoren der Gruppe im Plateau, d.h. das Rauschen wird unterdrückt, während echte Signale nicht abgeschnitten werden.

Bei der Messung hat sich herausgestellt, dass B5-1 kein Signal und T1-6 und T2-2 nur sehr schwache Signale liefern. Die drei Szintillatoren wurden bei der Wahl der Schwelle nicht berücksichtigt. Die auf diese Weise bestimmten Plateau-Grenzwerte sind zusammen mit den eingestellten Diskriminatorschwellen für die neun Szintillator-Gruppen in Tabelle 6.1 zu finden. Falls die Kalibrierung zu einem späteren Zeitpunkt wiederholt werden soll, ist die Logik im FPGA und der ECL-NIM-Wandler vorhanden. Die Verbindungen von den Trefis zum Wandler müssen jedoch nach Bedarf gesteckt werden. Außerdem müssen die Kabel der einzelnen Szintillatoren an unabhängige Kanäle der Trefis angeschlossen werden. Zur Konfiguration der Messung liegt ein Skript vor, mit dem ein

Wert für die PDL eingestellt werden kann. Außerdem muss ausgewählt werden, welche Szintillatorebene verzögert werden soll. Mit dem Skript werden die zehn Zähler für Koinzidenzen automatisch gestartet, nach einer wählbaren Zeit gestoppt und anschließend automatisch ausgelesen.

Gruppe	T1	T2	T3	T4	B1	B2	B3	B4	B5
Maximale Minimalschwelle (mV)	120	130	150	150	100	105	100	75	60
Minimale Maximalschwelle (mV)	175	200	300	200	135	120	120	85	85
Gewählte Schwelle (mV)	150	165	225	175	120	110	110	80	75

Tab. 6.1.: Diskriminatorschwellen für die einzelnen Gruppen. Um die maximale Minimalschwelle bzw. minimale Maximalschwelle zu bestimmen, wurden die Plateaus der einzelnen Kurven miteinander verglichen.

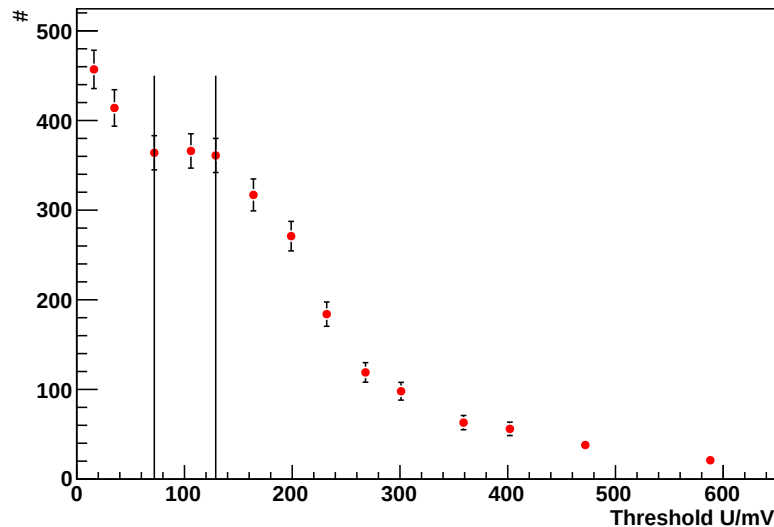


Abb. 6.9.: Die Koinzidenzen zwischen dem Szintillator B2-1 und der oberen Ebene bei verschiedenen Diskriminatorschwellen für B2-1. Die vertikalen Linien markieren das Plateau.

6.4.4. Zählraten der einzelnen Szintillatorgruppen

Mit den Signalen der Szintillatoren sollen Trigger für das Supermodul erzeugt werden. Um die Triggerraten für kosmische Teilchen abzuschätzen (vgl. Kap. 8.1) und um die Gruppen nach der Kalibrierung untereinander zu vergleichen, wurden die Zählraten der einzelnen Gruppen über zehn Minuten bestimmt.

In Tab. 6.2 bzw. Abb. 6.13 sind die Ergebnisse der neun Gruppen zusammengefasst. Die Zählraten der Gruppen einer Ebene sind vergleichbar. Schwankungen innerhalb einer Ebene lassen sich darauf zurückführen, dass die Szintillatoren in [Bat07] nach Qualität sortiert wurden. Die höhere Zählrate

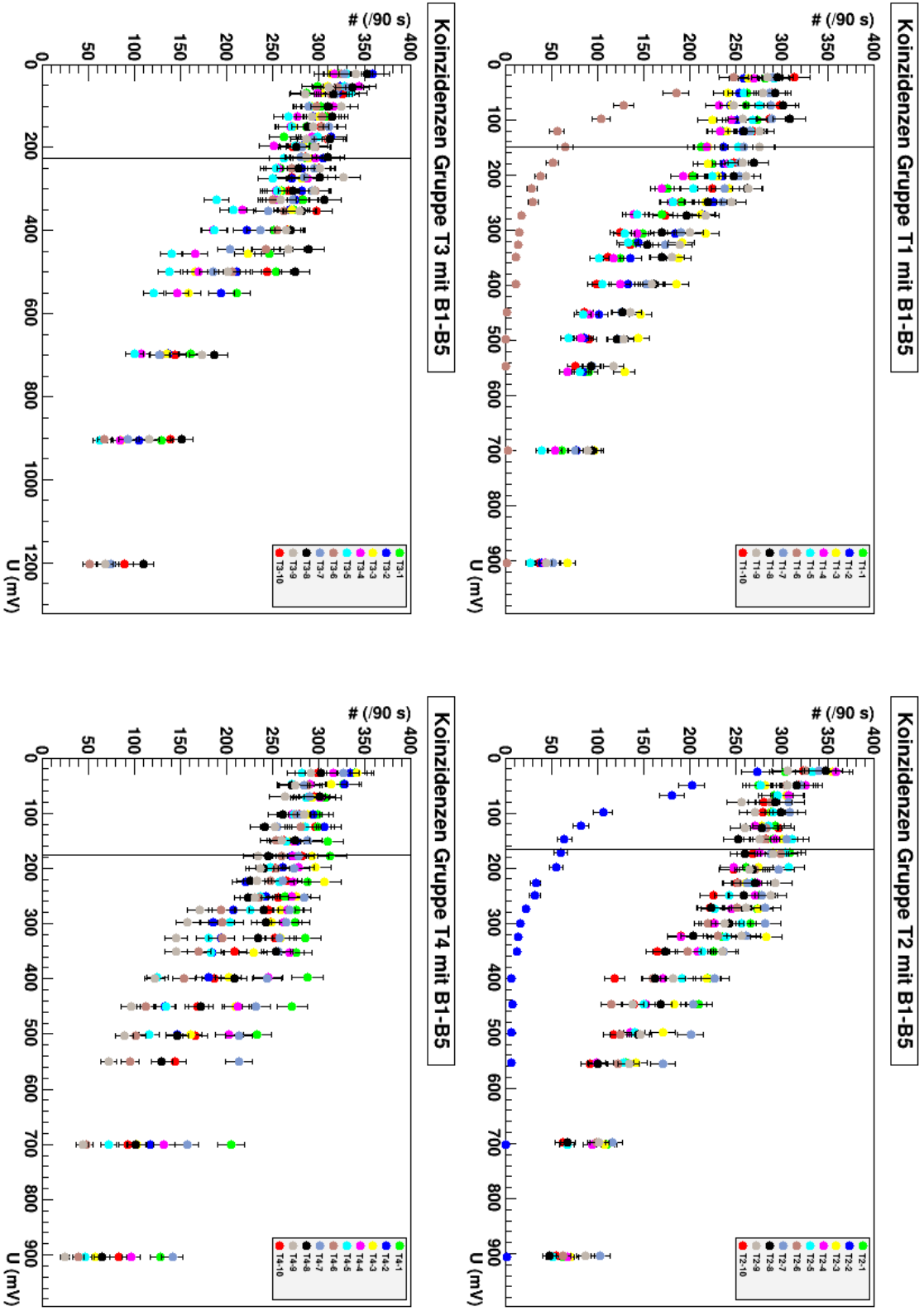


Abb. 6.10.: Die Koinzidenzen der Szintillatoren aus T1 (links oben), T2 (rechts oben), T3 (links unten) und T4 (rechts unten). Die vertikale Linie markiert die eingestellte Diskriminatorschwelle. T1-6 und T2-2 liefern nur sehr schwache Signale bzw. ist kein Plateau zu beobachten. Sie wurden bei der Wahl der Schwelle ignoriert.

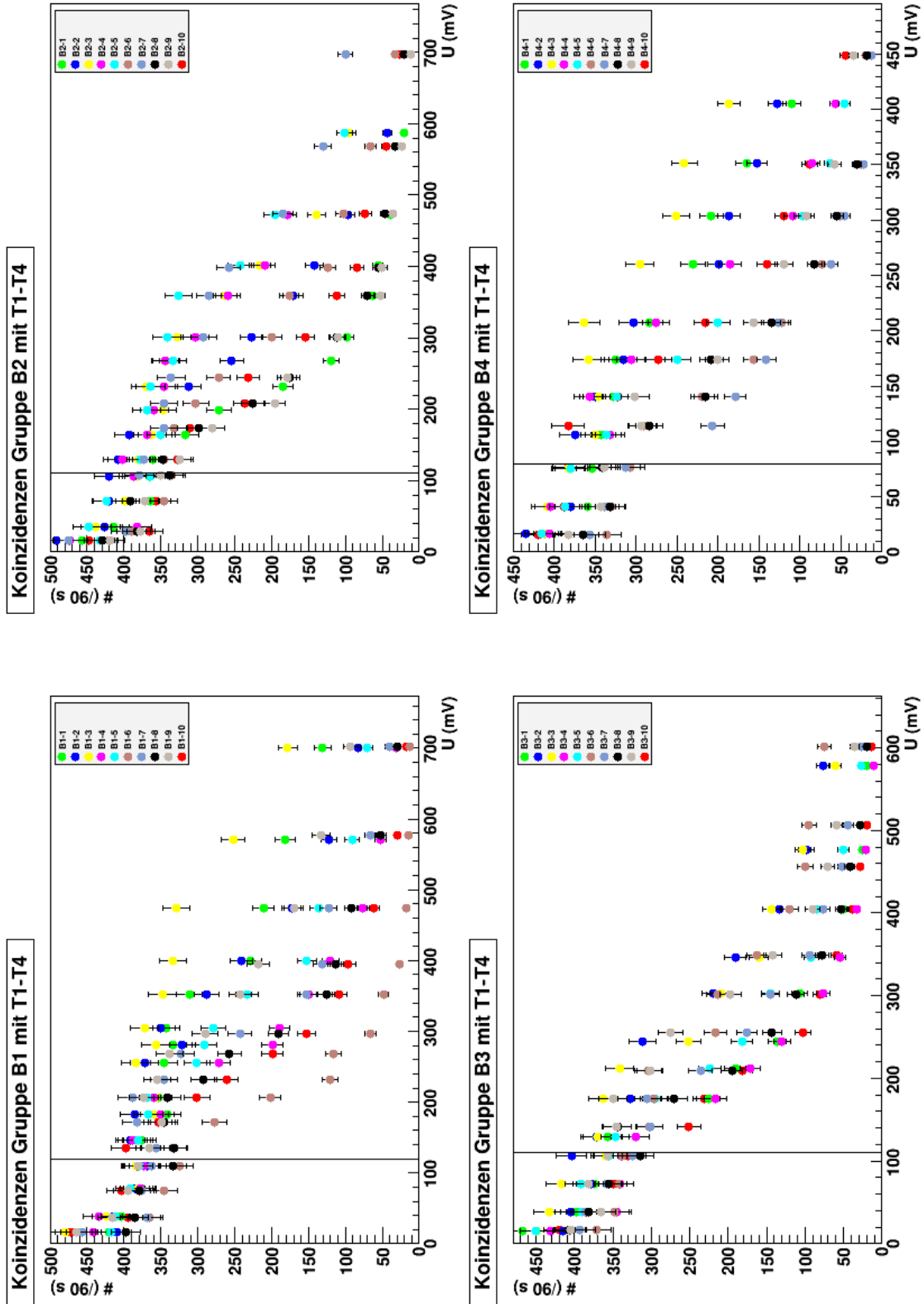


Abb. 6.11.: Die Koinzidenzen der Szintillatoren aus B1 (links oben), B2 (rechts oben), B3 (links unten) und B4 (rechts unten). Die vertikale Linie markiert die eingestellte Diskriminatorschwelle.

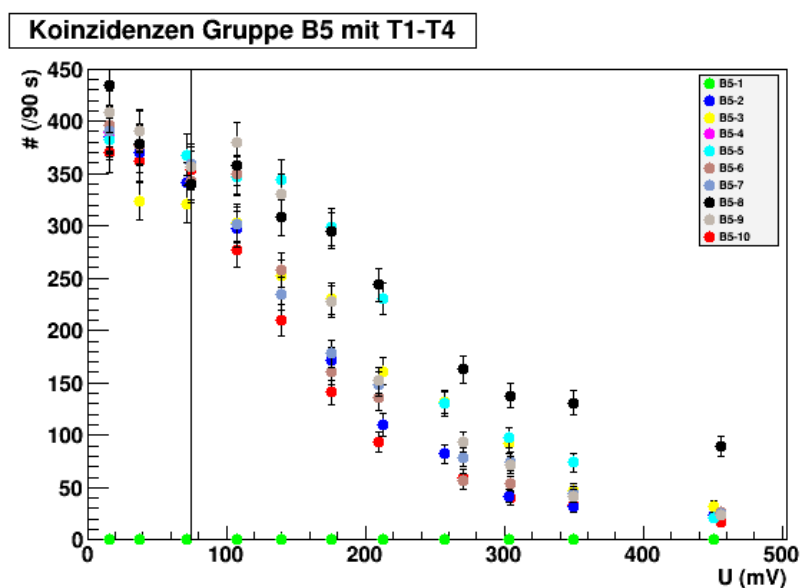


Abb. 6.12.: Die Koinzidenzen der Szintillatoren aus B5. Die vertikale Linie markiert die eingestellte Diskriminatorschwelle. B5-1 liefert kein Signal und wurde für die Wahl der Schwelle ignoriert.

von Top im Vergleich zu Bottom lässt sich auf die größere aktive Fläche einer Gruppe aus Top im Vergleich zu einer Bottom-Gruppe zurückführen.

Gruppe	T1	T2	T3	T4	B1	B2	B3	B4	B5
Signale (s^{-1})	337	400	361	378	234	237	200	225	200

Tab. 6.2.: Zählraten der einzelnen Gruppen. Die Fehler sind alle <1 , sind aber nur für große Statistiken gültig.

In Tab. 6.3 bzw. Abb. 6.14 sind die Zählraten der beiden gesamten Ebenen, die in der Triggereinheit gezählt wurden, gezeigt. Daneben ist jeweils die summierte Zählrate der einzelnen Gruppen zu sehen. Die Zählrate der Triggereinheit liegt etwas tiefer als die summierte Rate, da ein Ereignis, bei dem zwei oder mehr Szintillatoren aktiv sind, nur einfach gezählt wird.

Daraus lässt sich ablesen, dass die Anzahl der Ereignisse mit mehr als einem aktiven Szintillator bei Top ≤ 18 Ereignisse pro Sekunde bzw. bei Bottom ≤ 11 Ereignisse pro Sekunde ist. Diese Ereignisse sind entweder Teilchenschauer oder zufällige doppelte Treffer. Außerdem ist es möglich, dass ein kosmisches Teilchen mit extrem flachen Einfallswinkel zwei Szintillatorengruppen einer Ebene trifft.

Im Folgenden wird nicht die Summe der einzelnen Gruppen, sondern die Zählrate der Triggereinheit verwendet, da nur mit dieser Rate Trigger erzeugt werden können.



Abb. 6.13.: Hier sind die Zählraten der einzelnen Szintillatorengruppen aufgetragen. Es ist gut zu erkennen, dass jeweils die Gruppen der unteren Ebene und der oberen Ebene ungefähr gleiche Raten liefern. Der Unterschied zwischen einer Gruppe aus Top und einer Gruppe aus Bottom ist jedoch auffällig. Eine Erklärung dafür ist die größere aktive Fläche einer Top-Gruppe.

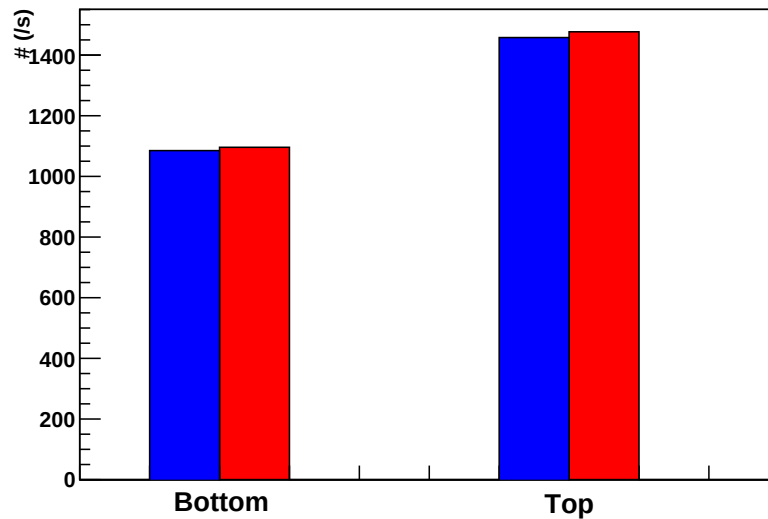


Abb. 6.14.: Zählraten von Top und Bottom. In blau ist die jeweilige Zählrate der Triggereinheit zu sehen. In Rot sind die summierten Zählraten der einzelnen Gruppen zu sehen. Diese ist höher, da der Zähler in der Triggereinheit ein Ereignis, bei dem zwei Sintillatoren aktiv sind, nur einfach zählt.

Gruppe	Top <i>Summe</i>	Top <i>Oder</i>	Bottom <i>Summe</i>	Bottom <i>Oder</i>
Signale (/s)	1476	1458	1096	1085

Tab. 6.3.: Zählraten der Szintillatorebenen. Mit *Summe* ist hier die Summe der Zählraten der einzelnen Gruppen bezeichnet. *Oder* gibt die Zählrate der Triggereinheit an. Die Fehler sind alle <1 , sind aber nur für große Statistiken gültig.

In diesem Kapitel wurde der Aufbau des Teststandes in Münster vorgestellt, mit dem kosmische Teilchen detektiert werden können, um einen Trigger für das Supermodul zu generieren. Dieser Aufbau wurde kalibriert. Dazu wurde die Logikeinheit aus [Bat07] durch eine VME-Karte mit einem frei programmierbaren FPGA ersetzt.

Die Kabellängen zwischen den Trefis und dem FPGA wurden für die Koinzidenzbedingung optimiert und die Diskriminatorschwellen der Trefis wurden unter Berücksichtigung aller Szintillatoren eingestellt. Damit ist die Kalibrierung der Szintillatoren bzw. ihrer Elektronik abgeschlossen. Im nächsten Kapitel wird die neue zentrale Triggereinheit vorgestellt.

7. Triggersystem

In diesem Kapitel wird das neue Triggersystem, das im Rahmen dieser Arbeit entwickelt wurde, vorgestellt. Zunächst wird eine Motivation gegeben, warum eine neue Triggereinheit, die die Signale der Szintillatoren und der GTU verarbeitet und Triggersequenzen erzeugt, sinnvoll ist. Danach werden der Aufbau und die einzelnen Bestandteile erläutert. Außerdem wird die Verzögerungszeit des Pretriggers abgeschätzt.

Im Anschluss wird die zentrale Zustandsmaschine vorgestellt, die im FPGA des V1495 die Triggerausgabe kontrolliert und Grundlage für alle Triggersequenzen (vgl. Kap. 8) ist.

Daraufhin wird ein interner Logikanalysator vorgestellt, der parallel zur Triggergeneration im FPGA läuft. Mit ihm können interne Signale überwacht und analysiert werden.

7.1. Motivation für eine neue Triggereinheit

Mit dem Koinzidenzaufbau aus [Bat07] konnten bereits kosmische Teilchen mit einem Supermodul aufgezeichnet werden. Dazu wurde ein Pretrigger generiert, wenn die Koinzidenzbedingung der Szintillatoren erfüllt war. Der Level 1 wurde von der GTU erzeugt.

Die Spuren von kosmischen Teilchen konnten jedoch nicht über einen langen Zeitraum aufgezeichnet werden, da in unregelmäßigen Abständen die Stromversorgung für die TRAPs zusammengebrochen ist. Mit der neuen Triggereinheit soll die Stabilität über einen langen Zeitraum gewährleistet werden.

Des Weiteren soll die neue Einheit einen erweiterten Funktionsumfang bieten, indem sie verschiedene Trigger für kosmische Teilchen und Triggersequenzen zum Testen kombiniert. Außerdem soll sie vollständig über ein TCP/IP-Netzwerk kontrolliert werden können, um die Benutzerfreundlichkeit zu verbessern.

Stabilität

Mit dem Aufbau aus [Bat07] ist es häufiger zu Abstürzen der TRAPs bzw. zum Zusammenbruch der Stromversorgung gekommen, wenn über einen langen Zeitraum Daten genommen wurden. Im Folgenden soll erläutert werden, wo die Ursachen dafür liegen.

Ein Supermodul braucht verschiedene Versorgungsspannungen. Neben den Hochspannungen für die Driftkammern, werden noch diverse Niederspannungen für die MCMs benötigt.

In einem Supermodul sind mehr als 5000 MCMs mit jeweils vier 120 MHz CPUs integriert, die eine Spannung von 1,8 V benötigen. Sind die CPUs eingeschaltet, braucht einer von ihnen zwischen 50–100 mA [Ang05]. Würden alle CPUs eines Supermoduls gleichzeitig mit maximalem Stromverbrauch aktiviert, würde eine Strom von ca. 1500 A fließen. Deswegen werden die TRAPs und insbesondere die Tracklet Prozessoren nur dann aktiviert, wenn sie tatsächlich benötigt werden.

Während des Betriebs des Supermoduls müssen also nur kurze Stromspitzen bei der Versorgungsspannung der CPUs ausgeglichen werden. Dazu werden während der Standby-Phasen Kondensatoren aufgeladen, die bei Bedarf entladen werden. Dadurch kann eine Stromversorgung verwendet werden, die sehr viel schwächer ist als eine, die für den Dauerbetrieb der TRAPs ausgelegt ist. Die Stromversorgung in Münster ist in drei unabhängige Einheiten unterteilt, die jeweils zwei Lagen des Supermoduls versorgen. Sie sind so konfiguriert, dass sie sich automatisch abschalten, wenn für die Versorgungsspannung der Tracklet Prozessoren (1,8 V) Ströme $I \geq 160$ A fließen. Das entspricht einem Maximalstrom von 480 A für das gesamte Supermodul.

Die MCMs erwarten im normalen Betrieb drei Triggerpulse (PT, L0 und L1)¹. Da alle Triggerpulse identisch sind und sich nur durch das Zeitfenster, in dem sie von den TRAPs akzeptiert werden, unterscheiden, können sie unter Umständen verwechselt werden. Werden die Pulse vom Supermodul anders interpretiert als vom Triggersystem erwartet, arbeiten Triggersystem und Supermodul nicht mehr synchron und die TRAPs können durch unerwartete Trigger abstürzen.

Stürzen die TRAPs in einem Zustand ab, in dem die Tracklet Prozessoren aktiv sind, werden diese nicht mehr deaktiviert und die Stromversorgung bricht zusammen.

Ein weiteres Szenario, das berücksichtigt werden muss, ist, dass die Triggerrate kurzzeitig zu hoch ist. Werden die Tracklet Prozessoren zu oft pro Zeiteinheit aktiviert, kann der Strombedarf nicht mehr gedeckt werden und die Versorgung bricht zusammen.

Die Stromversorgung ist zwar für jeweils zwei Lagen eines Supermoduls unabhängig, aber da die Triggersignale an alle Lagen gleichzeitig geschickt werden, bricht die Versorgung in der Regel trotzdem für das gesamte Supermodul zusammen. Nach einem Zusammenbruch können keine Daten mehr aufgezeichnet werden, bis die Stromversorgung wiederhergestellt und das Supermodul neu konfiguriert wird.

Die Supermodule stehen nur für einen begrenzten Zeitraum (in der Regel ein bis zwei Wochen) zur Datennahme zur Verfügung, so dass die Zeit möglichst effektiv genutzt werden muss, um eine ausreichende Statistik für die Kalibrierung zu erreichen. Kosmische Teilchen werden in der Regel über Nacht aufgenommen, wenn das Supermodul nicht für andere Tests benutzt wird. Da die Datennahme nicht überwacht wird, geht die Zeit zwischen einem Zusammenbruch und dem nächsten Morgen verloren.

Um einen stabilen Betrieb zu ermöglichen, muss also sicher gestellt werden, dass:

- das Triggersystem und die TRAPs synchron laufen.

¹ Mit dem Aufbau aus [Bat07] wurden nur PT und L1 erzeugt. Die TRAPs wurden dementsprechend konfiguriert, dass sie keinen L0 erwarten.

- die Triggerrate nie (also auch nicht kurzzeitig) zu hoch ist.

Um diese Anforderungen zu erfüllen, basiert die neue Triggereinheit auf einer Zustandsmaschine (vgl. Kap. 7.3).

Erweiterter Funktionsumfang und Benutzerfreundlichkeit

Zusätzlich zu den Triggern, die zur Aufzeichnung von kosmischen Teilchen dienen, werden während der Montage der Supermodule verschiedene Tests durchgeführt. Dafür werden diverse Trigger benötigt. Diese Testtrigger wurden bisher mit einer Elektronik erzeugt, die sehr unflexibel ist. Die Triggersequenzen sollen mit der neuen Triggereinheit erzeugt und verbessert werden.

Um zwischen den verschiedenen Triggern zu wechseln, mussten bisher immer Kabel umgesteckt werden. Mit der neuen Einheit soll es möglich sein, alle Einstellungen und insbesondere die Auswahl des Triggers mit einem Computer über ein TCP/IP-Netzwerk zu steuern. Dadurch wird es unter anderem möglich, die Triggereinheit mit Skripten oder Programmen, die bestimmte Tests durchführen, zu kontrollieren.

Um die Funktion der Triggereinheit zu überwachen, sollen außerdem Triggerraten, ausgewählter Trigger und weitere Informationen laufend ausgelesen und angezeigt werden (vgl. Kap. 9).

7.2. Aufbau in der Montagehalle in Münster

In diesem Unterkapitel wird der experimentelle Aufbau des Triggersystems in Münster vorgestellt. Dazu werden die einzelnen Komponenten und ihre Funktion beschrieben. Im Anschluss wird noch die Signallaufzeit des Pretriggers abgeschätzt.

In Abb. 7.1 ist der komplette Aufbau skizziert und in Abb. 7.2 ist ein Foto des elektrischen Aufbaus mit dem V1495, den Pegelwandler, TTCex und TTCvi zu sehen. Das V1495 sammelt die Signale der Szintillatoren und der GTU und generiert Triggersequenzen. Diese werden über optische Fasern an das Supermodul bzw. die GTU verteilt. Die GTU stellt einen L1-Beitrag (L1ctb kurz für *Level 1 Contribution*) zur Verfügung, sammelt die Daten des Supermoduls und gibt sie zur dauerhaften Speicherung an das DAQ-System weiter.

V1495

Das V1495 ist die zentrale Triggereinheit. Es verarbeitet die Signale der Szintillatoren. Zusätzlich berücksichtigt es gegebenenfalls das GTU-busy Signal und den Level-1-Beitrag der GTU.

Mit dem V1495 werden Triggersequenzen erzeugt, die aus Pretrigger, Level 0 und Level 1 bestehen. Das Supermodul erwartet diese drei Pulse jeweils mit einer Länge von 25 ns (vgl. Kap. 4.5).

Die GTU erwartet eigentlich ein Standard-TTC-Signal mit L0A (25 ns), L1A (50 ns), L1M und L2A bzw. L2R. Da die L1-Nachricht und der L2A nur spezielle Ereignis-Informationen des CTP

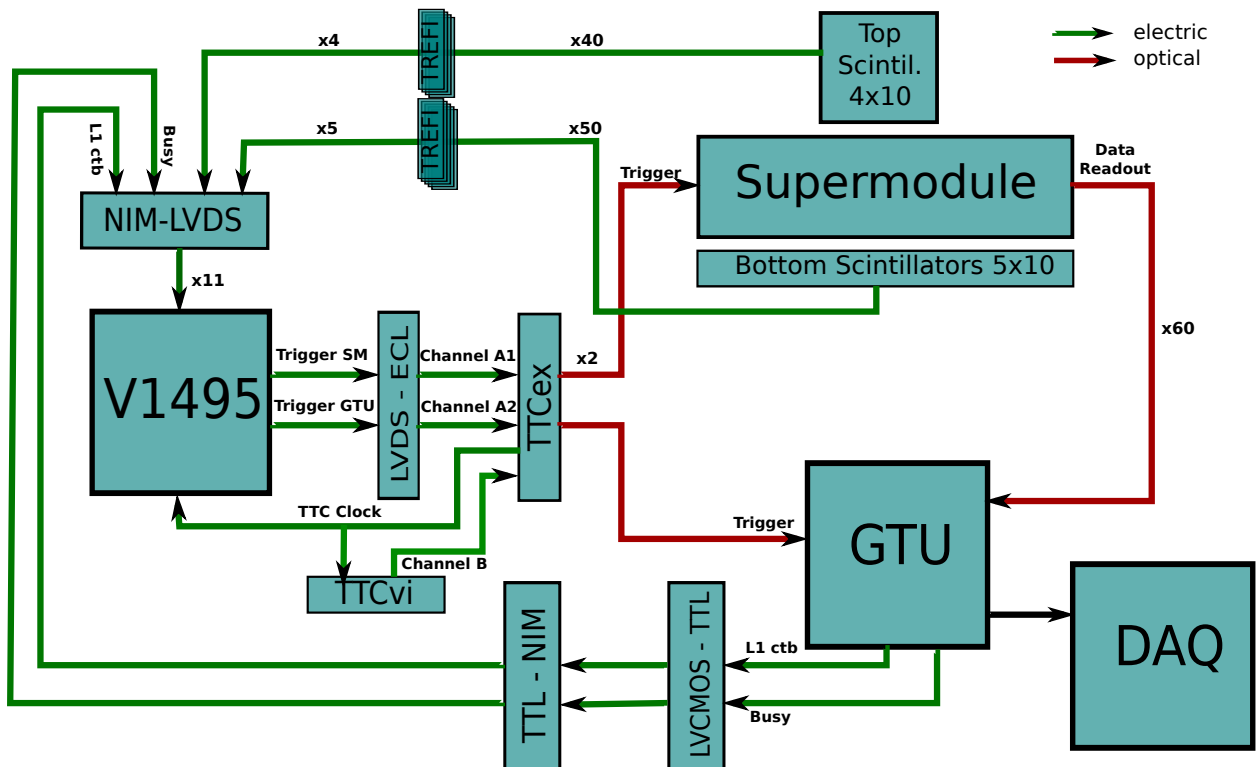


Abb. 7.1.: Hier ist der schematische Aufbau des Triggersystems gezeigt. Das V1495 verarbeitet die Signale der Szintillatoren und das Busy-Signal und den Level 1-Beitrag der GTU. Es gibt die Triggersignale über separate Kanäle an die GTU und Supermodul weiter. Die GTU sammelt die Daten des Supermoduls und gibt sie an das DAQ-System weiter. Die Zahlen neben den Verbindungen geben die Anzahl der Kanäle an.

und des LHC enthalten (vgl. 3.3.4), die in Münster nicht zur Verfügung stehen, wird auf diese verzichtet.

Da das Supermodul und die GTU verschiedene Trigger benötigen, werden sie über getrennte Kanäle des V1495 ausgegeben. Alle Pulse werden über den A-Kanal des TTC-Systems übertragen.

TTCex – Laser Transmitter

Der TTCex wird in zwei unabhängigen Hälften betrieben. Das bedeutet, dass er zwei A/B-Kanal-Eingänge hat, die unabhängig voneinander in ein optisches Signal gewandelt werden. Dadurch ist es möglich, die GTU und das Supermodul mit nur einem TTCex mit verschiedenen Signalen zu versorgen.

Für das Signal für die GTU muss ein optischer Dämpfer (20 dB) benutzt werden, da ein Ausgang des TTCex normalerweise mehrere DCS-Boards versorgt. Dadurch ist die Leistung des optischen Signals für ein DCS-Board zu hoch und die Empfängerdiode des TTCrx könnte zerstört werden.

Damit das TTCex korrekt arbeitet, müssen die eingehenden Signale synchron zur TTC-CLK sein. Da der FPGA des V1495 auch mit der TTC-CLK getaktet ist, muss nur die Phase angepasst

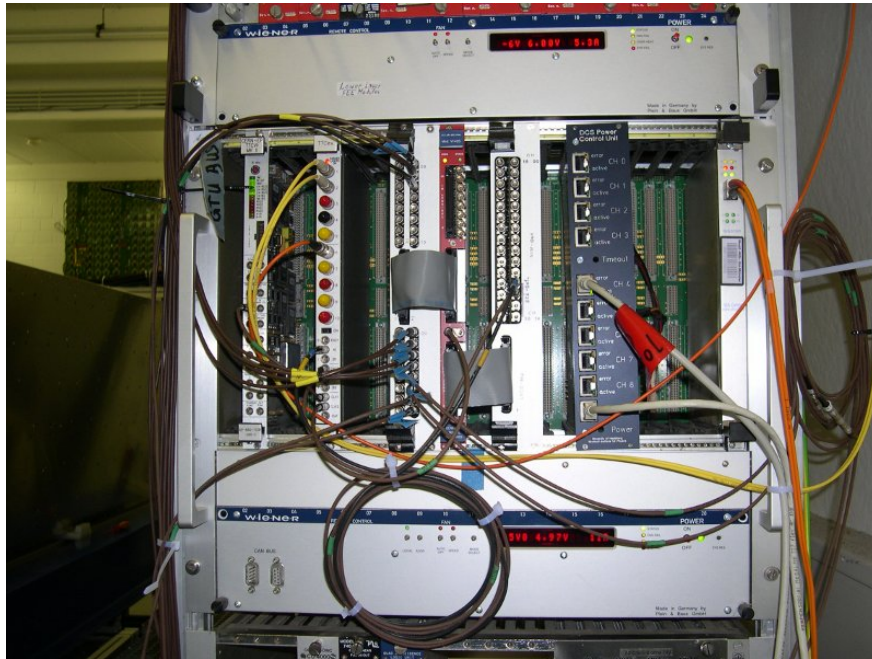


Abb. 7.2.: Von links nach rechts: TTCvi, TTCex, NIM-LVDS, V1495, LVDS-NIM. Rechts daneben befinden sich noch weitere Module, die aber nicht für das Triggersystem benutzt werden.

werden. In Abb. 7.3 ist dargestellt, welche Anforderungen an ein externes Signal gestellt sind, damit es ordnungsgemäß vom TTCex verarbeitet wird. Dabei muss $\Delta s \geq 5,1 \text{ ns}$ und $\Delta h \geq 3,4 \text{ ns}$ sein.

Zur Einstellung wurden die Eingangssignale vom TTCex zusammen mit der TTC-CLK auf einem Oszilloskop dargestellt und die Triggerpulse so verzögert, dass die richtige Phasenbeziehung gegeben ist. Die Kabel zwischen V1495 und TTCex entsprechen jetzt einer Verzögerung von 10 ns.

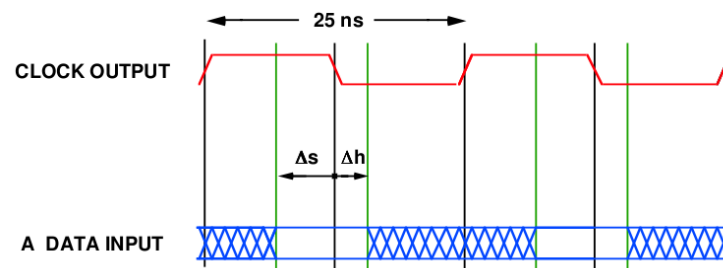


Abb. 7.3.: Synchronisation der Triggersignale am TTCex-Eingang. Die Signale müssen auf die fallende Flanke der TTC-CLK synchronisiert werden. Dabei muss $\Delta s \geq 5,1 \text{ ns}$ und $\Delta h \geq 3,4 \text{ ns}$ sein. Die Bereiche, die mit X markiert sind, werden vom TTCex ignoriert. Data Input ist in diesem Fall das Triggersignal vom V1495. Bild aus [Tay]

TTCvi

Über den TTCvi wird der B-Kanal bereitgestellt. Dieser wird nur dazu benötigt, damit die DCS-Boards bzw. TTCrx auf den Modulen und in der GTU die TTC-CLK wiederherstellen können. Das Signal wird aufgeteilt und an den B1- und B2-Eingang angeschlossen. Darüber hinaus wird der B-Kanal bzw. das TTCvi nicht benötigt.

Global Tracking Unit

Die Global Tracking Unit erfüllt die gleiche Aufgabe wie am CERN. Nachdem es einen L0 von der Triggereinheit empfangen hat, ist es bereit, die Spursegmente des Supermoduls zu empfangen. Nach deren Auswertung gibt es einen L1-Beitrag an die Triggereinheit zurück.

Erhält die GTU einen L1A von der Triggereinheit, bereitet es sich auf den Empfang der Rohdaten vor. Im Normalfall würde die GTU nach dem Empfang auf einen L2A bzw. L2R von dem CTP warten. In Münster sind jedoch keine anderen Detektoren vorhanden, die einen L2R verlangen könnten. Deshalb wird innerhalb der GTU nach $88\ \mu\text{s}$ automatisch ein L2A erzeugt und die Daten werden vom DAQ-System dauerhaft gespeichert.

Die GTU wird für die verschiedenen Trigger unterschiedlich oder gar nicht benutzt. Das genaue Verhalten wird an entsprechender Stelle in Kap. 8 erläutert.

7.2.1. Signallaufzeit des Pretriggers

Der Pretrigger braucht durch die verwendete Elektronik und Kabel eine gewisse Zeit bis er bei den TRAPs ankommt. Damit diese Signallaufzeit ausgeglichen werden kann, kann die Ausgabe der Analog-Digital-Wandler um bis zu sechs Timebins, bzw. 600 ns verzögert werden (vgl. Kap. 4.6.2). Dadurch können alle Datenpunkte der Ionisationsspur aufgezeichnet werden, obwohl sie bei Ankunft des Pretriggers bereits teilweise in der Vergangenheit liegen.

Um die passende Verzögerung zu bestimmen, wurden die Signallaufzeiten abgeschätzt. Sie sind in Tabelle 7.1 zusammengefasst. Aus [Bat07] sind die mittleren Verzögerungen durch die Szintillatoren, Photoervielfacher und die Trefis bekannt. Die Kabellängen bzw. die Laufzeiten der Signale von den Trefis zum V1495 wurden bereits für die Koinzidenzbedingung (vgl. Kap. 6.3) bestimmt.

Die Verarbeitungszeit im V1495 und die Signallaufzeit bis zum DCS-Board am Supermodul wurden mit einem Oszilloskop bestimmt, indem die Signale der zu vermessenden Verzögerung mit gleich langen Kabeln an das Oszilloskop angeschlossen und verglichen wurden.

Die Verarbeitungszeit im V1495 variiert aufgrund der Synchronisation der Eingangssignale um ca. 20 ns. Für die Berechnung wird der Mittelwert von 100 ns zugrunde gelegt.

Die Gesamtverzögerung liegt bei 321 ns. Je nachdem welcher Triggermodus aktiv ist und wieviel Logik benötigt wird, um den Pretrigger zu generieren, können weitere Verzögerungen von bis zu 50 ns (2 Takte) auftreten. Für die Datennahme werden die ADCs auf die maximale Verzögerung von 600 ns eingestellt. Dadurch sind vor jeder Ionisationsspur zwei bis drei leere Timebins vorhanden,

	Verzögerung (ns)
Szintillator	3
Photovervielfacher	20
Trefi	14
Trefi–V1495	64
V1495	100
V1495–Supermodul (DCS)	120
Summe	321

Tab. 7.1.: Verzögerungszeiten der einzelnen Elemente.

die zur Bestimmung der Nulllinie der ADCs genutzt werden können.

7.3. Zustandsmaschine

In diesem Unterkapitel wird die globale Zustandsmaschine (FSM) beschrieben, mit der die Triggersequenzen in der neuen Triggereinheit generiert werden. Indem eine FSM benutzt wird, soll gewährleistet werden, dass die Triggereinheit und die TRAPs immer synchron laufen. Aus diesem Grund orientiert sich die FSM auch stark an der FSM der TRAPs, die zunächst vorgestellt wird. Im Anschluss wird die FSM der Triggereinheit vorgestellt und die Zeitvorgaben diskutiert.

7.3.1. Zustandsmaschine der TRAPs

In Abb. 7.4 ist die Zustandsmaschine der TRAPs schematisch dargestellt. Ist das Supermodul für die Datennahme konfiguriert, sind die TRAPs im Zustand *Wait for Pretrigger*. Wird ein Pretrigger empfangen, werden in den Zuständen *Preprocessing* und *Tracklet Processing* die Rohdaten im Zwischenspeicher abgelegt und die Spursegmente berechnet. Kommt kein Level 0 im vorgesehenen Zeitfenster, werden die Spursegmente und die Rohdaten in *Clear* wieder verworfen und die Elektronik wartet auf den nächsten Pretrigger.

Wenn ein Level 0 geschickt wird, werden die Spursegmente zur GTU übertragen (*Tracklet Transmission*). Anschließend warten die TRAPs in *Wait for Level-1 Accept* auf den Level 1–Trigger. Wenn im vorhergesehenen Zeitfenster kein L1A ankommt, werden die Rohdaten in *Clear* verworfen und der TRAP ist bereit für eine neue Triggersequenz. Bei einem L1A werden die Rohdaten an die GTU geschickt. Anschließend warten die TRAPs auf einen neuen Pretrigger.

Die Wartezeiten zwischen Pretrigger und Level 0 und zwischen Pretrigger und Level 1 können bei der Konfiguration der MCMs eingestellt werden und müssen mit den Zeitvorgaben der Triggereinheit übereinstimmen, damit ein funktionierender und stabiler Betrieb gewährleistet ist (siehe Tab. 7.2).

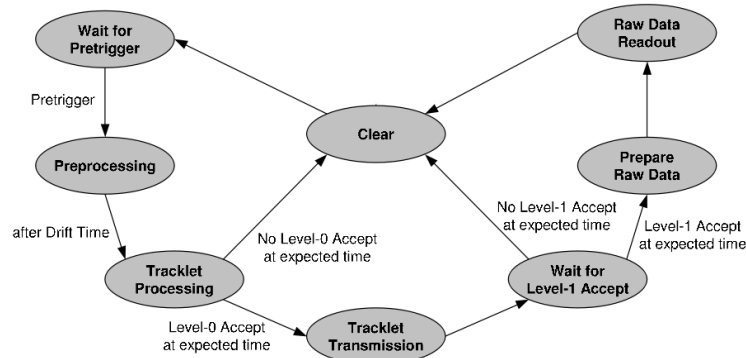


Abb. 7.4.: Die FSM wie sie in den TRAPs implementiert ist. Die TRAPs erwarten eine Triggersequenz mit Pretrigger, Level 0 und Level 1. Wird kein L0 bzw. L1 in den jeweiligen Zeitfenstern geschickt, werden die Rohdaten verworfen und die TRAPs warten nach einer Wartezeit auf einen neuen Pretrigger. Nach einem L1A werden die Rohdaten an die GTU übertragen. Danach kann eine neue Triggersequenz gestartet werden. Bild aus [Ang05]

7.3.2. Zustandsmaschine der Triggereinheit

In Abb. 7.5 ist die Zustandsmaschine der Triggereinheit dargestellt. Sie stellt das Grundgerüst für jede Triggersequenz dar. Dadurch ist die Reihenfolge der Triggerpulse festgelegt und es sind starre Zeitvorgaben zwischen den Pulsen definiert.

Ob ein Pretrigger, Level 0 oder Level 1 zum gegebenen Zeitpunkt erzeugt wird, hängt vom gewählten Triggertyp ab. Generell wird die Entscheidung, ob ein Trigger ausgegeben wird oder nicht, immer in der Triggereinheit getroffen. Welche Bedingungen dafür im Einzelnen erfüllt sein müssen, ist für jeden Trigger unterschiedlich. Für einige werden nur interne Signale berücksichtigt, während z.B. die Trigger für kosmische Teilchen die Szintillatoren und den Level 1–Beitrag der GTU berücksichtigen. Eine ausführliche Beschreibung der verfügbaren Trigger ist in Kap. 8 zu finden.

Nach einer vollständigen Sequenz wechselt die FSM in *Readout*, wo sie darauf wartet, dass die Rohdaten zur GTU geschickt werden. Bei bestimmten Triggersequenzen zum Testen werden die Daten nicht gespeichert und gar nicht von der GTU verarbeitet. In diesem Fall wird das GTU–busy nicht beachtet und es kann ein wählbarer Zeitraum eingestellt werden, in dem die FSM im Status *Readout* bleibt, bevor sie in *Clear* wechselt.

Nach einer abgeschlossenen oder abgebrochenen Trigger–Sequenz wechselt die FSM in den Zustand *Clear*, der neue Triggersequenzen für einen gewissen Zeitraum blockiert, damit die TRAPs sich auf den nächsten Pretrigger vorbereiten können.

7.3.3. Zeitvorgaben der Zustandsmaschinen

In Tabelle 7.2 sind die Zeitvorgaben der FSM der Triggereinheit bzw. der TRAPs zusammengefasst. Die Triggereinheit ist mit der TTC–CLK (40,079 MHz) getaktet, während die TTC–CLK für die

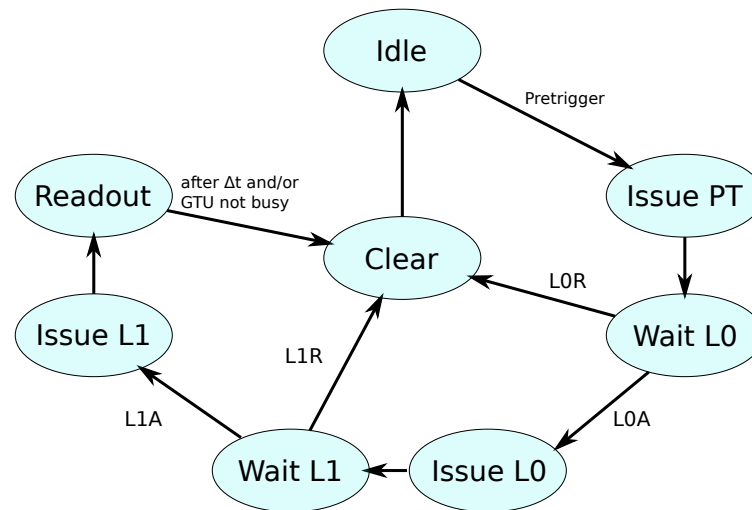


Abb. 7.5.: Die FSM der Triggereinheit. Für die verschiedenen möglichen Trigger müssen unterschiedliche Bedingungen erfüllt sein, damit eine Triggersequenz mit einem Pretrigger gestartet wird, bzw. mit einem L0 oder L1 fortgesetzt wird. L0R bzw. L1R stehen hier für ausbleibende L0- bzw. L1-Pulse. Nach einem L1A wechselt die FSM in *Readout*. In diesem Status wird entweder eine feste Zeit gewartet oder solange, bis die GTU nicht mehr beschäftigt ist. Falls die Sequenz abgebrochen oder abgeschlossen wird, geht die FSM in einen *Clear*-Status, in dem sie darauf wartet, dass die TRAPs für den nächsten Pretrigger bereit sind.

FSM der TRAPs mit 3 dividiert wird. Das entspricht einem 120,237 MHz Taktgeber.

Da in Münster außer dem Supermodul keine anderen Detektoren vorhanden sind, sind die Zeitvorgaben im Vergleich zum Betrieb am CERN großzügiger gewählt. Das hat den Vorteil, dass z.B. neue Versionen der GTU getestet werden können und für die Spurrekonstruktion mehr Zeit zur Verfügung steht. Dadurch kann die Funktionalität zunächst ohne straffe Zeitvorgaben geprüft und gegebenenfalls optimiert werden.

Die Zeitvorgabe PT – Clear ist nur wichtig für solche Trigger, die das Busy-Signal der GTU ignorieren. Dann wird ein ca. 490 μs langes Zeitfenster nach dem L1A erzeugt, in dem keine neuen Triggersequenzen gestartet werden dürfen. Dadurch wird verhindert, dass ein Pretrigger gesendet wird, wenn die TRAPs die Rohdaten versenden. Diese Option wurde insbesondere für den Stresstest, der die Datennahme mit hohen Raten simuliert, eingebaut und wird noch in Kap. 8.2.2 diskutiert.

Die TRAPs brauchen ca. 500 ns [Ang10] nach einem L0R, L1R bzw. nach dem Auslesen der Rohdaten, um für die nächste Sequenz bereit zu sein. Hier wurde deshalb konservativ 1 μs Wartezeit gewählt, da dadurch die Triggerrate nur minimal verringert wird, während ein stabiler Betrieb gewährleistet ist.

Die Zeiten können über Register² verändert werden. Es ist jedoch zu beachten, dass die Zeiten an die verfügbaren Konfigurationen der TRAPs und die Konfiguration der GTU angepasst sind.

² Eine Liste aller verfügbaren Register mit ihren Adressen ist in Anhang E zu finden. In der Regel erfolgt der Zugriff auf die Register jedoch über die Benutzerschnittstelle, die in Kap. 9 vorgestellt wird.

7. Triggersystem

Falls Änderungen an der FSM vorgenommen werden, müssen diese Konfigurationen auch angepasst werden.

Seitdem die Zustandsmaschine der Triggereinheit mit diesen Zeitvorgaben zur Triggenergeneration eingesetzt wird, ist es nicht mehr zu Abstürzen der TRAPs bzw. zum Zusammenbruch der Stromversorgung gekommen. Die einzelnen Trigger werden in Kap. 8 vorgestellt.

	Trigger		TRAP		
	TTC-CLK Takte	Zeit (μs)	TRAP Takte	Zeit (μs)	Δt (μs)
PT – L0	46	1,12	138	1,12	$\pm 0,06$
PT – L1	316	7,88	949	7,89	$\pm 0,06$
PT – Clear	20.000	499	–	–	–
Clear	40	1	–	–	–

Tab. 7.2.: Die Zeitvorgaben der FSMs. Die TRAPs akzeptieren L0 und L1 in einem 120 ns breiten Zeitfenster Δt . Die Zeit PT-Clear und Clear sind für die TRAPs nicht definiert. Sie richten sich danach, wann die TRAPs die Rohdaten verschickt haben bzw. für einen neuen Trigger bereit sind.

7.4. Interner Logikanalysator

In diesem Kapitel soll ein *Interner LogikAnalysator* (ILA) vorgestellt werden, der im Rahmen dieser Arbeit programmiert wurde, um Fehler während der Entwicklung und dem Gebrauch der Triggereinheit zu identifizieren. Mit ihm können Triggersequenzen und andere Signale aufgezeichnet und analysiert werden.

Ein Logikanalysator hat eine ähnliche Funktion wie ein Oszilloskop. Er stellt allerdings nur die logischen Pegel, also logisch 1 (*high*) oder logisch 0 (*low*) dar. Dafür ist er in der Lage, viele Signale gleichzeitig zu visualisieren. Außerdem können komplexe Trigger-Bedingungen benutzt werden, um den gewünschten Zeitbereich bzw. Effekt zu beobachten.

In Münster steht kein Logikanalysator zur Verfügung. Deshalb wurde eine Logik in den FPGA implementiert, die interne Signale überwacht und aufzeichnet. Der ILA besteht aus zwei Teilen, die im Folgenden beschrieben werden:

1. Logik im FPGA des V1495, die parallel zur Logik der Triggereinheit läuft.
2. Ein C++-Programm, das den ILA steuert, ausliest und die Daten auf zwei Arten visualisiert.

7.4.1. Logik im FPGA

Mit der Logik wird bestimmt, welche Signale der Triggereinheit aufgezeichnet werden sollen und unter welchen Triggerbedingungen die Aufzeichnung gestartet bzw. gestoppt wird. Um den Verlauf der Signale zu speichern, werden sie mit jedem TTC-Takt in einen RAM-Block geschrieben. Der Schreibvorgang wird durch ein Triggersignal gestoppt, woraufhin der RAM ausgelesen werden kann. Der Status des ILA wird durch eine Zustandsmaschine gesteuert.

RAM

Die Signale, die überwacht werden sollen, werden jeden Takt im RAM gespeichert und sollen über den VME-Bus ausgelesen werden. Im FPGA sind M4K-RAM-Blöcke vorgesehen. Das ist ein Speicher, der auf verschiedene Weisen genutzt werden kann (für weitere Informationen sei auf [Alt08] verwiesen). Für den ILA ist ein 2-Port RAM-Block gut geeignet.

In Abb. 7.6 ist ein schematisches Bild des implementierten RAM mit den Ein- und Ausgängen gezeigt. Er hat einen Port für Schreibzugriffe (*data*) und einen für Lesezugriffe (*out*). Die aktuelle Lese- bzw. Schreibadresse wird dementsprechend in *rdaddress* bzw. *wraddress* gespeichert. Für den Schreibport gibt es zusätzlich noch das Bit *wren*, mit dem der Schreibzugriff erlaubt bzw. verboten wird.

Alle Signale sind mit der TTC-CLK getaktet und ein Schreibzugriff erfolgt immer mit der fallenden Flanke der TTC-CLK. Der gesamte Block besteht aus vier M4K-Blöcken und hat eine Gesamtkapazität von 16384 Bit bzw. 1024 Datenworten mit 16 Bit. Das bedeutet, dass maximal 16 Signale über einen Zeitraum von 1024 TTC-Takten (ca. 25 μ s) aufgezeichnet werden können.

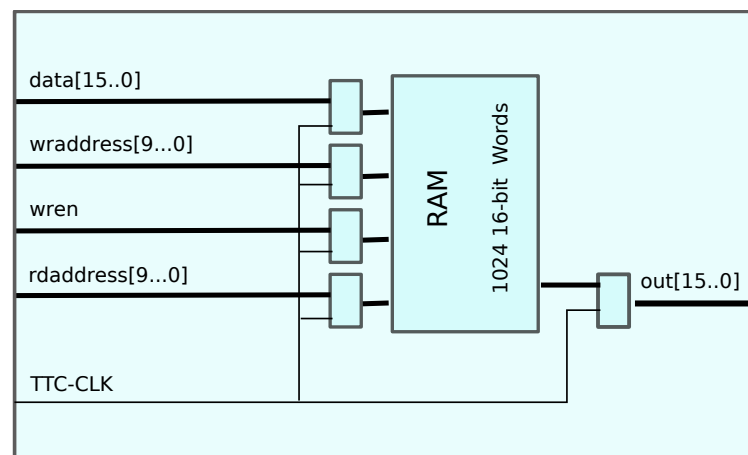


Abb. 7.6.: Schematische Darstellung des 2-Port RAM im FPGA. Es kann gleichzeitig gelesen und geschrieben werden, da für Lese- bzw. Schreiboperationen separate Ports vorgesehen sind. Der Block kann 1024 Datenworte mit einer Breite von 16 Bit abspeichern.

Das ist für die Fehlersuche und die Aufnahme von Triggersequenzen vollkommen ausreichend, wenn geeignete Triggerbedingungen vorhanden sind, die die Aufnahme starten bzw. stoppen.

Die 16 Signale, die aufgezeichnet werden, werden ständig in den RAM geschrieben, wobei die Schreibadresse bei jedem TTC-Takt um eins erhöht wird. Wird die Schreibadresse 1023 erreicht, springt der Adressenzähler zurück auf Adresse 0. Dadurch sind immer die letzten 1024 Takte im Speicher enthalten.

Trigger zum Starten bzw. Stoppen des ILA

Der RAM wird ca. alle 25 μs neu überschrieben, wenn das *wren*-Bit gesetzt ist. Um interessante Ereignisse analysieren zu können, muss der Schreibvorgang im richtigen Moment gestoppt werden. Dazu sind Trigger, die eine oder mehrere Bedingungen überprüfen, gedacht. Unabhängig davon, welche Trigger verwendet werden, gibt es die beiden Modi *Pre* und *Post*:

- Pre bedeutet, dass *wren* sofort auf low gesetzt wird, wenn eine Triggerbedingung erfüllt ist. Dadurch werden keine neuen Daten in den RAM geschrieben. Es sind dann die letzten 1023 Takte vor dem Trigger gespeichert.
- Post bedeutet, dass noch einmal eine beliebige Anzahl von Takten (max. 1023) Daten in den RAM geschrieben werden und dann gestoppt wird.

Diese beiden Trigger-Modi sind unabhängig von den verschiedenen Triggern, die benutzt werden können. Die Trigger werden über eine Bit-Maske mit 16 möglichen Triggern ausgewählt. Dabei können einer oder mehrere Trigger gleichzeitig aktiv sein. Einige der möglichen Trigger sind:

- Ein Pretrigger, Level 0 oder Level 1 wurde ausgegeben.

- Fallende Flanke GTU–busy.
- Top Szintillator oder Bottom Szintillator aktiv.

Das sind Standard–Beispiele, die universell benutzt werden können. Für konkrete Probleme muss aber unter Umständen ein komplexerer Trigger implementiert werden, der auf das Problem angepasst ist. Das ist insbesondere der Fall, wenn seltene Ereignisse gesucht werden. Für den ILA wird auch eine FSM verwendet. Nachdem ein Trigger ausgewählt wurde, kann dieser vom Benutzer aktiviert werden (vgl. Kap. 9.2.3). Dann wechselt die FSM in einen Status, in dem die Triggerbedingung(en) überprüft werden. Ist eine Bedingung erfüllt, wird für die eingestellte Anzahl an Takten weiter in den RAM geschrieben und anschließend wechselt die FSM wieder in den Wartestatus. Zusätzlich wird ein Bit aktiviert, das dem Benutzer anzeigt, dass getriggert wurde. Sobald ein Trigger den ILA gestoppt hat, kann der Inhalt des RAM ausgelesen werden. Anschließend können wieder Trigger aktiviert und das nächste Ereignis aufgenommen werden.

7.4.2. Steuerprogramm des ILA

Der ILA wird über ein Steuerprogramm kontrolliert und ausgelesen. Es ist außerdem in der Lage, den Inhalt des RAM mit ASCII–Zeichen in einer Konsole oder über ein ROOT–Makro graphisch zu visualisieren.

Das Steuerprogramm greift über den VMEbus auf Register im FPGA zu, wodurch z.B. die Bit–Maske für die Trigger eingestellt, der Trigger(–Modus) ausgewählt oder der Trigger aktiviert werden kann.

Außerdem liest es den RAM aus. Dazu wird zunächst die letzte RAM–Adresse, in die geschrieben wurde, ausgelesen. Diese Adresse plus eins ist dann der älteste gespeicherte Takt. Jetzt werden die nächsten 1023 Adressen ausgelesen, wobei wieder von 1023 auf 0 gesprungen wird. Die Vektoren, die so ausgelesen werden, werden verarbeitet und dargestellt.

In Abb. 7.7 ist ein kleiner Auszug aus einer Ausgabe auf der Konsole zu sehen. Hier wurde auf einen L1 getriggert. Die Signalkürzel sind in Tab. 7.3 zusammengefasst.

Alternativ können die Daten mit ROOT ausgewertet werden. Für die Auswertung ist ein Makro vorhanden, das automatisch ausgeführt wird und die Signale darstellt. Das hat den Vorteil, dass alle 1024 Takte gleichzeitig betrachtet werden können, und bei Bedarf die Vergrößerungsfunktion von ROOT benutzt werden kann. Ein Nachteil ist allerdings, dass der Status der FSM nicht dargestellt werden kann. In Abb. 7.8 ist eine Triggersequenz dargestellt.

Diese Bilder werden später zur Erläuterung der einzelnen Trigger benutzt, um die Sequenzen und das Zusammenspiel der Signale zu veranschaulichen.

In diesem Kapitel wurde die neue Triggereinheit, basierend auf einem FPGA, vorgestellt. Mit

Konsole	ROOT	Beschreibung
SM	SM Out	Trigger, die an das Supermodul ausgegeben werden
GTU	GTU Out	Trigger, die an die GTU ausgegeben werden
PT	Pretrigger	Pretrigger
L0	Level 0	Level 0
CT1	L1ctb in	Level 1–Beitrag von der GTU – Eingang
CT2	L1ctb syn	Einsynchronisierter Level 1–Beitrag der GTU
L1	Level 1	Level 1
BU1	Busy in	GTU–busy Signal – Eingang von der GTU
BU2	Busy used	GTU–busy intern mit Überschreiboption
TOP	SZ Top	Signal der oberen Szintillatorebene
BTM	SZ Btm	Signal der unteren Szintillatorebene
COI	Coinc	Koinzidenz zwischen Top und Bottom
L1O	L1 o-ride	Ob der Level 1 intern generiert wird oder nicht
STATE	–	Status der FSM

Tab. 7.3.: Beschreibung der einzelnen Signale, die vom ILA überwacht werden.

Takt	SM	GTU	PT	L0	CT1	CT2	L1	BU1	BU2	TOP	BTM	COI	L1O	STATE
0	.	.	.	.	.	.	.	X	.	.	.	.	.	wait_l1
1	.	.	.	.	.	.	.	X	.	.	.	.	.	wait_l1
2	.	.	.	.	.	.	.	X	.	.	.	.	.	wait_l1
3	.	.	.	.	.	.	.	X	.	.	.	.	.	wait_l1
4	X	X	.	.	.	.	X	X	.	.	.	.	.	send_l1
5	.	X	.	.	.	.	.	X	.	.	.	.	.	readout
6	.	.	.	.	.	.	.	X	.	.	.	.	.	readout
7	.	.	.	.	.	.	.	X	.	.	.	.	.	readout
8	.	.	.	.	.	.	.	X	.	.	.	.	.	readout
9	.	.	.	.	.	.	.	X	.	.	.	.	.	readout

Abb. 7.7.: Hier wurde auf einen Level 1 getriggert. Dabei wurde der ILA so eingestellt, dass er nach dem Level 1 noch 1020 Takte weiterschreibt. Dadurch sind die vier Takte vor dem Level 1 zu sehen. In diesem Fall wurde das GTU–busy intern überschrieben. Dadurch ist BU2 low, während BU1 high ist. Wenn ein L1 generiert wird, wird ein 25 ns Puls an das Supermodul und ein 50 ns Puls an die GTU ausgegeben. Rechts ist der Status der FSM zu beobachten.

einer Zustandsmaschine wird gewährleistet, dass die Triggereinheit synchron zum Supermodul und zur GTU läuft. Außerdem werden Wartezeiten zwischen den einzelnen Triggersequenzen definiert. Dadurch wird garantiert, dass die Stromversorgung für die Tracklet Prozessoren nicht zusammenbricht.

Des Weiteren wurde ein interner Logikanalysator vorgestellt, mit dem diverse Signale überwacht werden. Dadurch ist es möglich verschiedene Triggersequenzen zu analysieren oder Fehler zu finden.

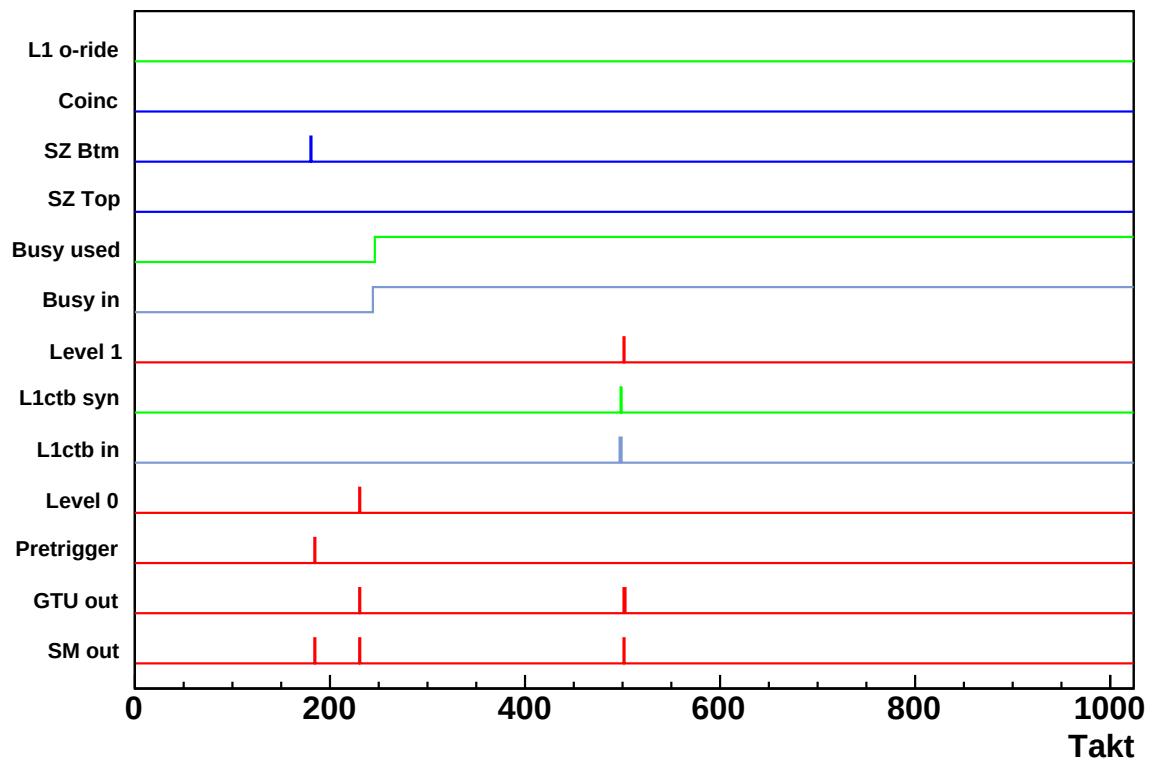


Abb. 7.8.: Ein Beispiel für eine Triggersequenz, aufgenommen mit dem internen Logikanalysator. Die roten Signale sind Trigger, die erzeugt bzw. an das Supermodul und die GTU ausgegeben werden. Die graublauen Signale sind die eingehenden Signale von dem GTU und die grünen sind interne Signale, die für die Triggersequenzen generiert werden. In dunkelblau sind die Signale der Top und Bottom Szintillatoren und deren Koinzidenz dargestellt. Hier ist eine Triggersequenz zu sehen, bei der keine Koinzidenz für einen Pretrigger erforderlich ist, sondern nur ein Treffer in Bottom. Stattdessen muss aber ein L1-Beitrag von der GTU kommen, damit ein L1 ausgegeben wird.

8. Trigger

In diesem Kapitel werden die verfügbaren Trigger vorgestellt. Dazu zählen zwei verschiedene Trigger, mit denen kosmische Teilchen aufgezeichnet werden können. Es wird gezeigt, wie die beiden Trigger durch den neuen Aufbau und die Triggereinheit kombiniert werden können.

Daraufhin werden einige Trigger vorgestellt, die zum Testen des Supermoduls verwendet werden. Darunter ist ein Stresstest, mit dem die Datennahme mit hohen Triggerraten simuliert werden kann und ein Trigger, mit dem das Rauschen der einzelnen Kanäle der MCMs aufgenommen werden kann.

Alle Trigger werden mit der FSM, die in Kap. 7.3.2 vorgestellt wurde, erzeugt. Dadurch wird garantiert, dass die Triggereinheit immer synchron zur Elektronik des Supermoduls und der GTU läuft. Außerdem werden Pausen zwischen den Triggersequenzen gefordert, die verhindern, dass die Triggerrate zu hoch wird. Wie ein Trigger ausgewählt wird, wird in Kap. 9.2.3 erläutert.

8.1. Trigger für kosmische Teilchen

In diesem Unterkapitel werden die Trigger diskutiert, die zur Verfügung stehen, um die Spuren von kosmischen Teilchen mit dem Supermodul aufzunehmen.

Dazu wird mit den Szintillatoren ein Pretrigger generiert, um die Elektronik des Supermoduls aufzuwecken. Um den Pretrigger zu generieren, werden zwei verschiedene Bedingungen genutzt:

- Koinzidenz zwischen einem Top und einem Bottom Szintillator
- Signal von einem Bottom Szintillator

Bei einer Koinzidenz wird automatisch die komplette Triggersequenz durchlaufen. Wenn jedoch nur ein Signal von Bottom für den Pretrigger gefordert wird, wird noch der Level 1-Beitrag der GTU, die eine Online-Spurrekonstruktion durchführt, ausgewertet. Die beiden Trigger werden im Folgenden vorgestellt.

8.1.1. Koinzidenztrigger

Bereits mit dem Aufbau aus [Bat07] konnte ein Pretrigger generiert werden, wenn in beiden Szintillatortoren gleichzeitig ein Signal erzeugt wurde. Diese Funktionalität wurde für die neue Triggereinheit übernommen und verbessert.

Im Folgenden wird die Logik für die Koinzidenz vorgestellt. Im Anschluss wird eine kurze Messung beschrieben, mit der die Pulslänge von Bottom und Top optimiert wurde. Zum Schluss wird die Performance des Koinzidenztriggers diskutiert.

Logik um Koinzidenzen zu detektieren

Zunächst wurde die Logik der Koinzidenzeinheit aus [Bat07] exakt nachgebildet, indem die Eingangssignale von Top und Bottom jeweils mit einem ODER zu einem Signal zusammengefasst und anschließend auf Gleichzeitigkeit geprüft wurden. Damit die Zeitinformation der Signale dabei nicht verloren geht, dürfen die Signale erst nach der Koinzidenz mit der TTC-CLK einsynchronisiert werden, da sie sonst um bis zu 25 ns verschoben werden können und echte Koinzidenzen unterdrückt werden. Da die Statistik für die Kalibrierung entscheidend ist, müssen möglichst viele kosmische Teilchen, die das Supermodul durchqueren, detektiert werden.

Um die Stabilität der Triggereinheit zu gewährleisten, ist es allerdings wünschenswert, dass alle Signale einsynchronisiert werden, da das Verhalten von asynchronen Signalen nur schwer vorher-sagbar ist. Außerdem sollen die Signale von Top bzw. Bottom möglichst kurz sein, damit zufällige Koinzidenzen vermieden werden.

Um alle Koinzidenzen zu detektieren, während nur synchrone Signale verwendet werden und die Pulslänge von Top und Bottom minimiert ist, müssen die Signale von Top und Bottom zunächst mit einer höheren Frequenz einsynchronisiert und verglichen werden. Dazu wurde die TTC-CLK mit einer Phasenregelschaltung (PLL) um den Faktor fünf dividiert, so dass zusätzlich noch ein 200 MHz-Taktgeber zur Verfügung steht. Werden die Szintillatorsignale mit diesem Taktgeber synchronisiert, verschiebt man die Signale um maximal 5 ns, was durchaus akzeptabel ist.

In Abb. 8.1 ist eine schematische Übersicht der Logik, die programmiert wurde, um einen Pretrigger mit der Koinzidenzbedingung zu generieren, dargestellt. Die Signale der neun Trefis werden mit einem 200 MHz-Taktgeber einsynchronisiert. Dazu wird das Signal für einen Takt auf high gesetzt, wenn eine steigende Flanke beobachtet wird.

Danach werden die Signale mit einer Maske verglichen, wodurch einzelne Szintillatorgruppen deaktiviert werden können. Das ist zum Beispiel sinnvoll, wenn nur mit einem Stapel des Supermoduls kosmische Teilchen aufgenommen werden sollen. Außerdem kann so eine Winkelbedingung an die Flugbahnen der kosmischen Teilchen gestellt werden. Für einen Test der Online- Spurrekonstruktion einer neuen GTU-Version im September 2010 wurde z.B. ein Trigger benötigt, der möglichst senkrechte Spuren detektiert. Dazu wurde eine Hälfte von Top mittig über einem Stapel platziert und alle anderen Gruppen von Top deaktiviert. In Bottom wurden alle Gruppen deaktiviert, die nicht unter dem entsprechenden Stapel waren.

Nach der Synchronisation mit 200 MHz sind die Signale von Top und Bottom genau einen Takt, also 5 ns lang. Da die Koinzidenzkurve durch verschiedene Effekte ausgeschmiert wird (vgl. 6.3), werden echte Koinzidenzen bei zu kurzen Signalen unterdrückt. Deshalb müssen die Signale von Top und Bottom verlängert werden, bevor sie verglichen werden. Die optimale Pulslänge wird im

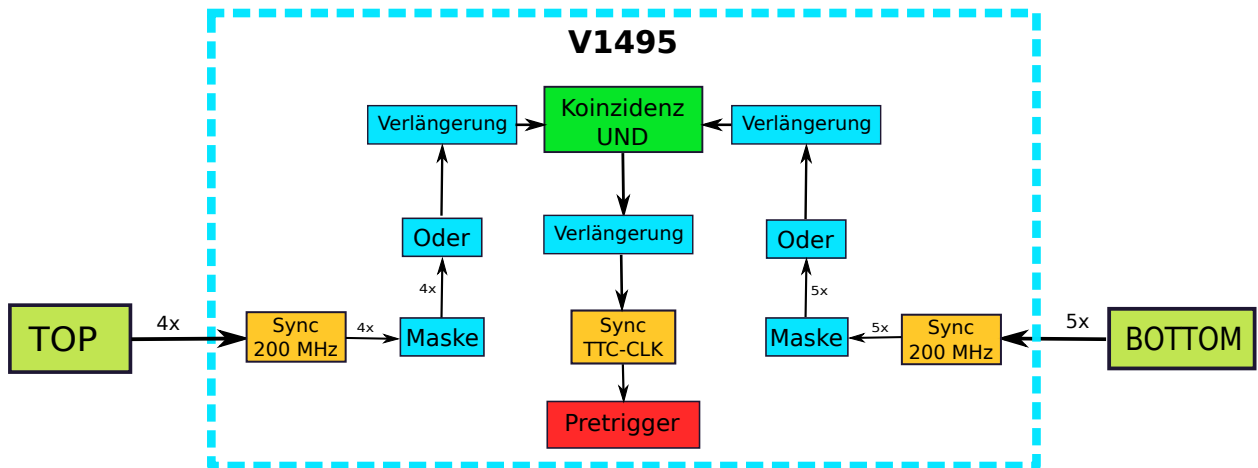


Abb. 8.1.: Die Logik um einen Pretrigger durch Koinzidenz zu erzeugen. Dazu werden die Signale von Top und Bottom zunächst mit 200 MHz einsynchronisiert. Mit einer Maske können beliebige Szintillatorgruppen deaktiviert werden. Bevor die Signale auf Koinzidenz geprüft werden, werden sie mit einem logischen ODER verschaltet und anschließend verlängert. Das Ausgangssignal der Koinzidenz wird ebenfalls verlängert, bevor es mit der TTC-CLK synchronisiert wird und einen Pretrigger generieren kann.

nächsten Unterkapitel bestimmt.

Nachdem die Signale verlängert wurden, werden sie mit einem UND auf Koinzidenz geprüft. Da sich die Signale möglicherweise nur einen 200 MHz-Takt überschneiden, muss das resultierende Signal ebenfalls verlängert werden, damit es korrekt mit der TTC-CLK synchronisiert und weiter verwendet werden kann. Hier wird es auf 50 ns gestreckt, damit gewährleistet ist, dass es mindestens einen vollen TTC-Takt high ist.

In Abb. 8.2 ist eine Triggersequenz des Koinzidenztriggers zu sehen, die mit dem ILA aufgenommen wurde. Ist der Koinzidenztrigger ausgewählt und die globale Zustandsmaschine im Status *idle*, wenn eine Koinzidenz gefunden wird, wird ein Pretrigger ausgegeben und die FSM gestartet. Level 0 und Level 1 werden zu den entsprechenden Zeitpunkten automatisch erzeugt.

Der Level 1-Beitrag der GTU wird für die Level 1-Entscheidung nicht berücksichtigt, da die Koinzidenzbedingung ausreichend ist. Außerdem erkennt die GTU in der Regel nur kosmische Teilchen, die durch einen und nicht mehrere Stapel im Supermodul fliegen (siehe Kap. 8.1.2 für Ausnahmen und weitere Informationen). Mit dem Koinzidenztrigger wird jedoch auch auf Teilchen getriggert, die einen flachen Einfallswinkel haben und mehr als einen Stapel durchqueren.

Das GTU-busy wird als Kriterium benutzt, um in der Zustandsmaschine von *Readout* auf *Clear* zu wechseln. Als zusätzliche Absicherung kann hier noch ein Zeitfenster benutzt werden, das verhindert, dass die FSM vorzeitig in *clear* wechselt. Dieses ist standardmäßig allerdings nicht aktiviert.

In Abb. 8.3 ist die gleiche Sequenz wie in 8.2 zu sehen, nur dass der Bereich des Pretriggers vergrößert wurde. Hier wurde eine Koinzidenz gefunden, obwohl sich die Signale von Top und Bottom nicht überschneiden.

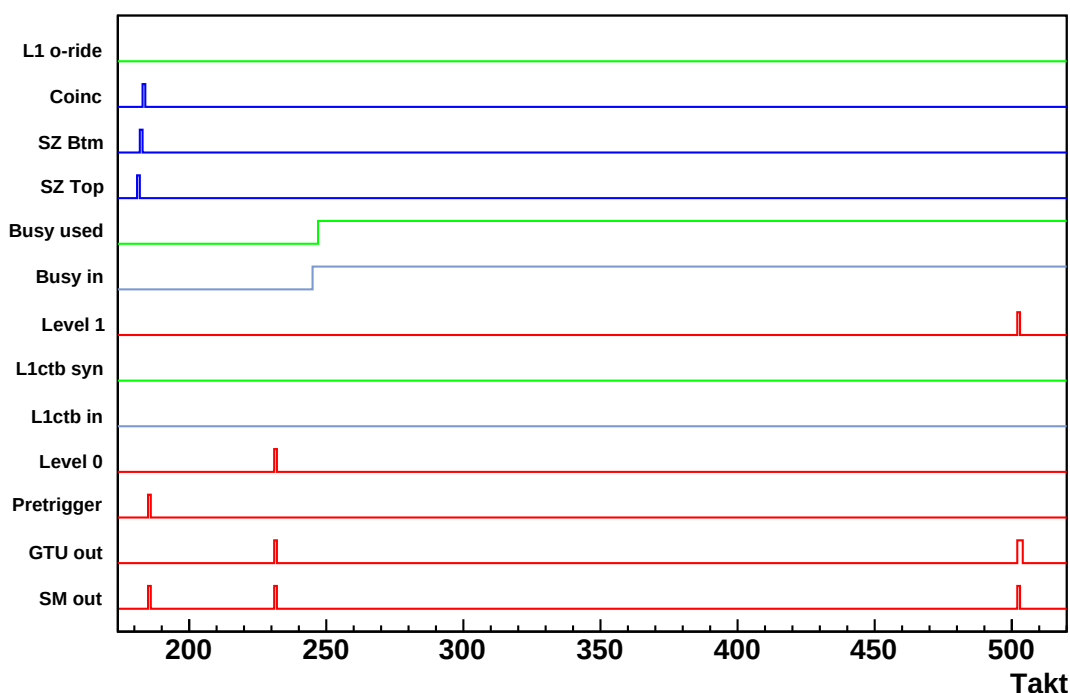


Abb. 8.2.: Eine Koinzidenz-Triggersequenz aufgenommen mit dem ILA. Ein Pretrigger wird durch eine Koinzidenz zwischen Top und Bottom erzeugt. Level 0 und Level 1 werden zum entsprechenden Zeitpunkt automatisch erzeugt. Die GTU wird mit dem L0 gestartet und gibt anschließend das GTU-busy Signal zurück. Ein L1-Beitrag der GTU wird nicht gefordert. In Abb. 8.3 ist eine Vergrößerung der Aufnahme um den Bereich des Pretriggers zu sehen.

Im ILA werden immer die Signale von Top und Bottom dargestellt, die synchron zur TTC-CLK sind. Das ist notwendig, da mit jedem TTC-Takt in den RAM geschrieben wird und Signale die synchron zum 200 MHz Taktgeber sind, nicht korrekt aufgezeichnet werden. In Abb. 8.3 ist zu sehen, dass Top und Bottom nicht im gleichen Takt high sind, aber trotzdem eine Koinzidenz gefunden wurde. Das liegt daran, dass sich die echten Signale von Top und Bottom in dem Grenzbereich zwischen zwei TTC-Takten überschneiden. Mit dem 200 MHz-Taktgeber wird dieser Bereich feiner abgetastet und die Koinzidenz wird gefunden.

Optimale Verlängerung für Top und Bottom

Um die optimale Verlängerung für die Top- und Bottom-Pulse zu bestimmen, wurde die Koinzidenzrate in Abhängigkeit von der Verlängerung der Pulse aufgenommen. Das Ergebnis ist in Abb. 8.4 zu sehen. Wie erwartet geht die Koinzidenzrate in Sättigung, wenn die Pulse lang genug sind, um die Flugzeitdifferenzen und die verschiedenen Verarbeitungszeiten der Elektronik auszugleichen (vgl. Kap. 6.3). Das ist hier bei 20 ns der Fall. Danach steigt sie nur noch minimal durch zufällige

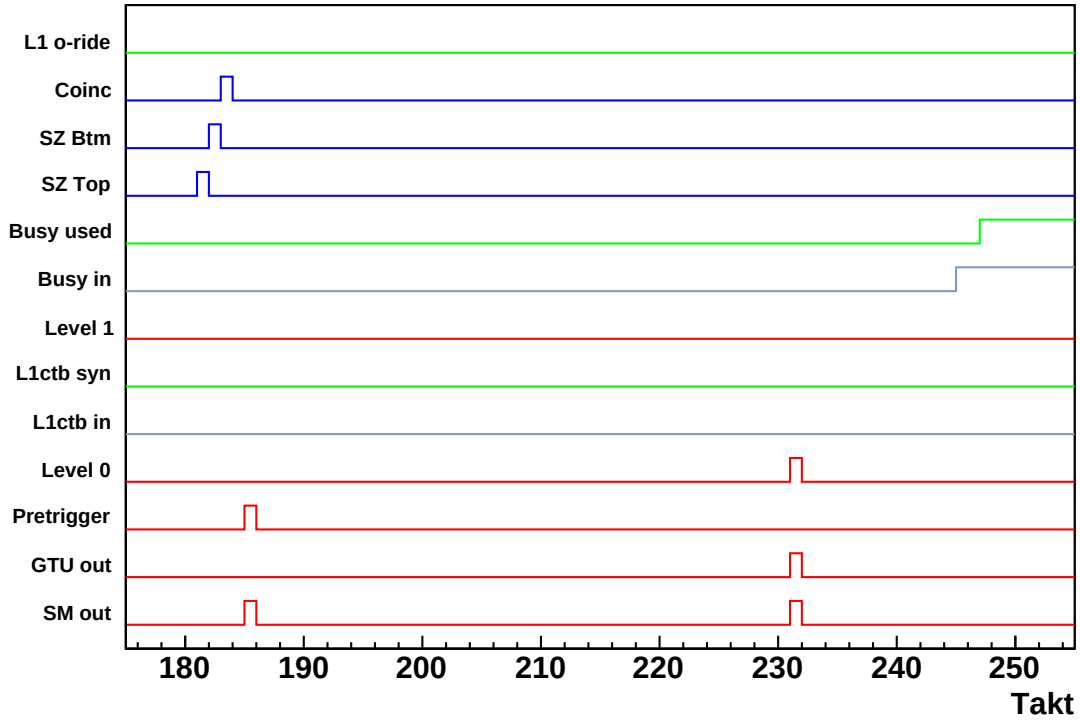


Abb. 8.3.: Es sind die Signale von Top und Bottom Signale gezeigt, die direkt mit der TTC-CLK einsynchronisiert wurden. Es ist zu sehen, dass sie in zwei aufeinanderfolgenden Takten high sind, obwohl eine echte Koinzidenz gegeben ist. In der Triggereinheit wird nur ein Pretrigger erzeugt, weil Top und Bottom mit dem 200 MHz Taktgeber abgetastet und verglichen werden. Außerdem ist zu beobachten, dass die Koinzidenz um einen TTC-Takt verzögert ist. Das liegt an der vergleichsweise aufwändigen Logik für die Koinzidenz.

Koinzidenzen. Dieser Effekt ist jedoch so klein, dass er im Diagramm nicht beobachtet werden kann.

Bottom liefert eine Rate von ca. 1090 Signalen pro Sekunde und Top ca. 1460 Signalen pro Sekunde (vgl. Kap. 6.4.4). Bei einer Signallänge von 25 ns ist die Rate der zufälligen Koinzidenzen K_{Zufall} verschwindend gering:

$$K_{Zufall} = 1090 \text{ s}^{-1} \cdot 1460 \text{ s}^{-1} \cdot 25 \cdot 10^{-9} \text{ s} = 0,040 \text{ s}^{-1} \quad (8.1)$$

Dabei verhalten sich die Pulslänge und die Rate zufälliger Koinzidenzen linear, d. h. durch eine Verdopplung der Pulslänge wird die Rate zufälliger Koinzidenzen auch verdoppelt. Da die Rate echter Koinzidenzen bei ca. 100 Signalen pro Sekunde liegt, können die zufälligen vernachlässigt werden. Für die Pulslänge von Top und Bottom wird hier ein Wert von 25 ns voreingestellt, da so auf keinen Fall echte Koinzidenzen unterdrückt werden.

Als Alternative zum Abtasten mit 200 MHz könnten die Signale auch mit der TTC-CLK abgetastet werden und auf zwei oder drei Takte verlängert werden. Dadurch sollten auch alle Koinzidenzen

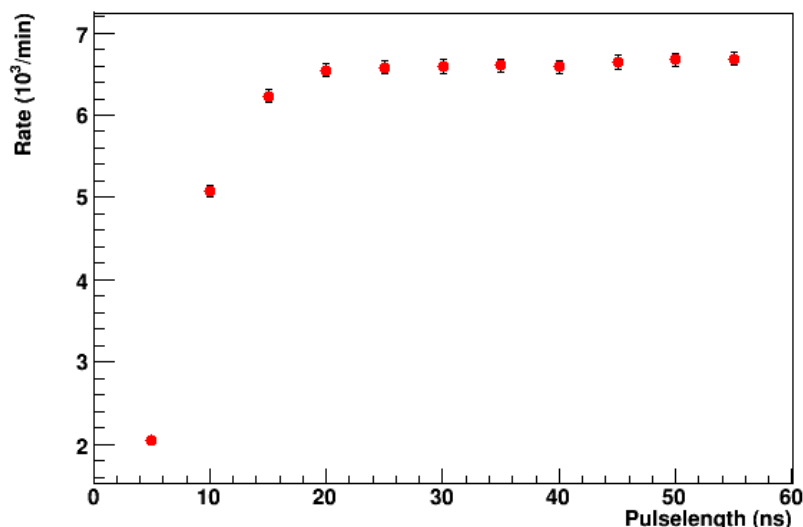


Abb. 8.4.: Die Konizidenzrate in Abhängigkeit der Pulslänge von Top und Bottom. Wie erwartet geht die Koinzidenz in Sättigung. Als Voreinstellung für die Länge wird 25 ns gewählt.

gefunden werden, wobei sich die Rate zufälliger Koinzidenzen erhöht. Diese wäre zwar immer noch zu vernachlässigen, jedoch um den Faktor zwei oder drei höher.

Aus diesem Grund und da die zwei integrierten PLLs des V1495 nicht für anderen Funktionen gebraucht werden, wird weiterhin mit 200 MHz einsynchronisiert.

Performance des Koinzidenztriggers

Wird Top mittig über Bottom positioniert, wird eine Koinzidenzrate von (108 ± 1) Signalen pro Sekunde erreicht. Um diesen Wert zu berechnen, wurden die Koinzidenzen von 10 Minuten aufsummiert. Bei dem Fehler muss berücksichtigt werden, dass die kosmische Strahlung Schwankungen unterliegt und er nur für große Statistiken gültig ist.

Da keine weiteren Beiträge einfließen, die die Triggersequenz gegebenenfalls abbrechen würden, entspricht die gemessene Rate der Koinzidenz auch ungefähr der aufgezeichneten Ereignisrate. Sie wird minimal reduziert, da eine Koinzidenz, die nicht im Zustand *Idle* stattfindet, verworfen wird. Bei einer Koinzidenzrate von 108 Signalen pro Sekunde und ca. 250 μs , die das GTU-busy pro Ereignis high ist, ergibt sich eine Totzeit von 2,7 %. Dadurch werden im Mittel drei Koinzidenzen pro Sekunde unterdrückt. Das entspricht einer Triggerrate von ca. 105 Signalen pro Sekunde.

Mit dem Aufbau aus [Bat07] wurde eine Triggerrate von ca. 90 Signalen pro Sekunde erreicht. Durch die neue Triggereinheit wurde diese Rate um mehr als 15% angehoben. Gründe dafür sind geringere Totzeiten der Triggereinheit, die sich jetzt streng an den Totzeiten des Supermoduls orientieren, und die verbesserte Kalibrierung der Diskriminatorschwellen.

Außerdem können jetzt einzelne Szintillatorgruppen deaktiviert werden, ohne dass Kabel um-

gesteckt werden müssen. Der größte Vorteil besteht jedoch darin, dass die Stromversorgung nicht mehr zusammengebrochen ist, seitdem die Zustandsmaschine für den Koinzidenztrigger benutzt wird (Juni 2010). Dadurch können Daten über einen langen Zeitraum aufgenommen werden.

8.1.2. GTU–Trigger

Zusätzlich zum Koinzidenztrigger gibt es einen GTU–Trigger, bei dem die Spursegment–Daten von der GTU online ausgewertet werden und ein Triggerbeitrag bereitgestellt wird. Mit dem Aufbau aus [Bat07] wurde der Beitrag der GTU direkt als Level 1 an das Supermodul weitergegeben. Mit dem neuen Triggersystem wird der Beitrag in der Triggereinheit verarbeitet und gegebenenfalls ein Level 1 an die GTU und das Supermodul ausgegeben.

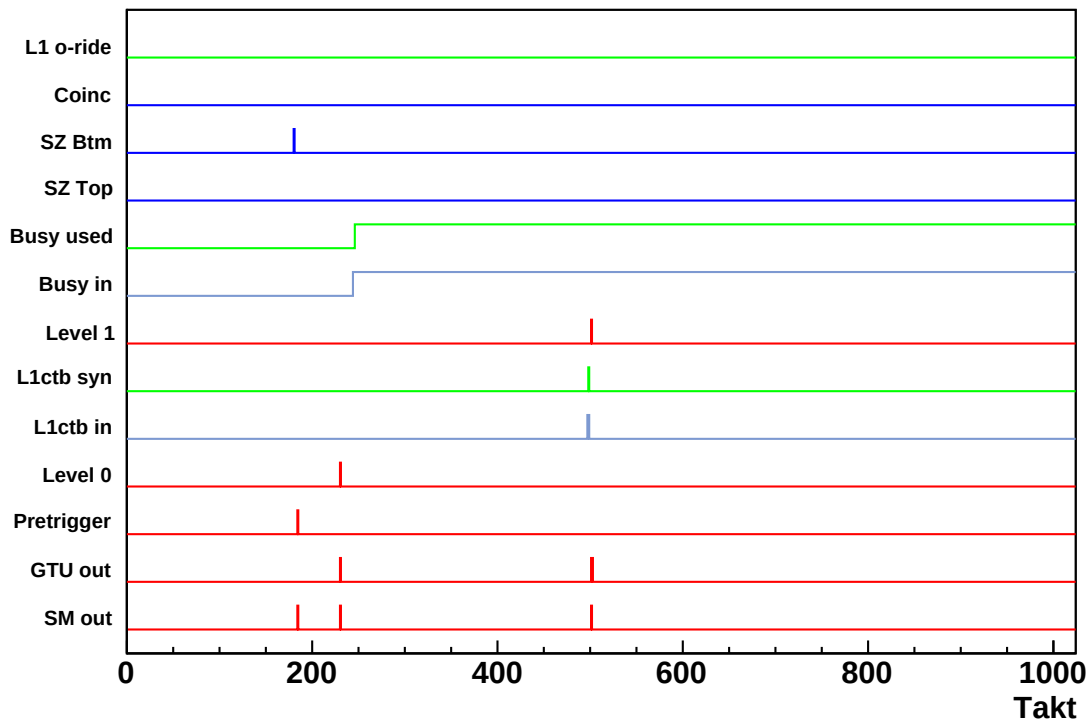


Abb. 8.5.: Hier ist eine normale Triggersequenz für den GTU–Trigger dargestellt. Durch ein Signal der unteren Szintillatorebene wird ein Pretrigger erzeugt. Der Level 0 wird automatisch zum entsprechenden Zeitpunkt erzeugt. Danach öffnet sich das Zeitfenster, in dem die L1ctb von der GTU akzeptiert wird. In diesem Bild kommt sie rechtzeitig und ein Level 1 wird erzeugt.

In Abb. 8.5 ist eine Aufnahme einer Triggersequenz des GTU Triggers mit dem ILA zu sehen. Beim GTU–Trigger wird nicht gefordert, dass eine Koinzidenz gefunden wird, um einen Pretrigger zu erzeugen. Da die Spursegmente von der GTU ausgewertet und so Spuren identifiziert werden, wird bereits ein Pretrigger ausgegeben, wenn Bottom ein Signal liefert. Die Rate von Bottom liegt

bei ca. 1090 Signalen pro Sekunde.

Aufgrund der geringeren Logik, ist dieser Pretrigger zwei Takte schneller als der Koinzidenzprettrigger. Damit die Signalformen der ausgelesenen Daten des Supermoduls für beide Trigger gleich sind, werden die Signale von Bottom um zwei TTC-Takte verzögert. Dadurch wird der Pretrigger immer zum gleichen Zeitpunkt relativ zum eigentlichen Ereignis erzeugt.

Der Level 0 wird wie bei dem Koinzidenztrigger automatisch im vorhergesehenen Zeitfenster erzeugt. Danach wartet die Triggereinheit auf einen Triggerbeitrag der GTU. Die L1ctb besteht aus einem 50 ns langen Puls, falls eine Spur gefunden wurde und keinem Puls, falls keine Spur gefunden wurde. Erhält die Triggereinheit keine L1ctb im vorhergesehenen Zeitfenster, wechselt die FSM in *Clear* und wartet anschließend in *Idle* auf den nächsten Pretrigger.

Kommt eine L1ctb im entsprechenden Zeitfenster (vgl. 7.3.3), wird ein Level 1 an das Supermodul und die GTU ausgegeben. Anschließend wartet die Triggereinheit darauf, dass die GTU nicht mehr beschäftigt ist, wechselt in *Clear* und wartet dann in *Idle* auf einen neuen Pretrigger bzw. ein Signal von Bottom.

Genauso wie beim Koinzidenztrigger können hier einzelne Szintillatorgruppen ausmaskiert werden. Da die fünf Gruppen von Bottom ziemlich genau mit den Stapeln übereinstimmen, können so Trigger für einzelne Stapel unterdrückt werden.

Einstellungen der GTU

Die echte Online-Spurrekonstruktion der GTU ist für kosmische Teilchen ungeeignet, da sie Teilchen von einem zentralen Kollisionspunkt erwartet. Kosmische Teilchen hingegen haben eine homogene Ursprungsverteilung.

Daher stellt die GTU ein *Cosmic-Tracking* bereit, das keine echte Spurrekonstruktion durchführt, sondern die in den Modulen deponierten Ladungsmengen überprüft. Zur Konfiguration werden zwei Schwellen eingestellt:

- Minimale integrierte Ladung pro Modul
- Minimale Anzahl getroffener Lagen pro Stapel

Mit der minimalen Ladung wird bestimmt, ob die Ladung ausreicht, um das entsprechende Modul als getroffen zu werten. Wie die Schwelle eingestellt werden muss, hängt von der Konfiguration der TRAPs ab. Für die in Münster verwendeten Konfigurationen hat sich eine Schwelle von 16 bewährt. Die Werte, die hier verglichen werden, entsprechen keinen echten Ladungen, sondern haben willkürliche Einheiten. Sie gehen auf Berechnungen der TRAPs zurück, die die Ladungen verarbeiten.

Mit der minimalen Anzahl der getroffenen Lagen pro Stapel legt man fest, wieviele Module in einem Stapel getroffen werden müssen, um eine L1ctb auszugeben. In Münster werden hier typischerweise vier Lagen gefordert. Die Schwellen für die Ladung bzw. Anzahl der Module pro Stapel können durch Register angepasst werden, falls andere Schwellen gewünscht sind.

Dass die GTU die Ladungen für jeden Stapel einzeln überprüft und nicht mehrere Stapel für eine Spur berücksichtigt, ist in der Architektur der GTU begründet (vgl. 4.6.3). Für jeden Stapel ist eine TMU vorhanden, mit der das Cosmic-Tracking durchgeführt wird. Die SMU sammelt die Ergebnisse der TMUs und gibt dementsprechend einen positiven L1-Beitrag aus, wenn eine der TMUs eine Spur gefunden hat. Mit dem Cosmic-Tracking werden in der Regel also nur Spuren gefunden, die nur einen Stapel des Supermoduls durchqueren. Ausnahmen stellen hier Teilchenschauer oder Teilchen, die durch vier Module eines Stapels und durch ein oder zwei Module eines benachbarten Stapels fliegen.

Die GTU hat verschiedene (Test-)Möglichkeiten um eine L1ctb zu erzeugen. Die aktuell eingestellten Methoden werden in der *contrib mask* zusammengefasst. Damit die Datennahme mit dem GTU-Trigger korrekt funktioniert, ist es entscheidend, dass in der *contrib mask* ausschließlich *cosmics* eingestellt ist.

Der Zeitpunkt, zu dem die GTU die L1ctb ausgibt, kann eingestellt werden. Im Moment ist er so konfiguriert, dass die L1ctb vier TTC-Takte vor Ende des *Wait L1*-Status an der Triggereinheit ankommt.

Performance des GTU-Triggers

Bottom hat eine Rate von ca. 1090 Signalen pro Sekunde. Wenn bereits eine Triggersequenz läuft, werden die Signale ignoriert und es wird kein neuer Pretrigger ausgegeben. Bei einer Triggerrate von 230 Signalen pro Sekunde und einer durchschnittlichen Auslesezeit von 250 μ s sinkt die Pretrigger-Rate durch die Totzeit auf ca. 1030 Signale pro Sekunde.

Der Level 0 wird automatisch erzeugt, womit seine Rate der Pretriggerrate entspricht. Ein Level 1 wird nur erzeugt, wenn eine L1ctb von der GTU kommt. Mit den oben genannten Einstellungen reduziert sich die Rate der L1ctb bzw. L1 auf ca. 230 Signale pro Sekunde. Damit ist die Rate des GTU-Triggers ca. um den Faktor zwei höher als die Rate des Koinzidenztriggers. Es wird also eine doppelt so hohe Statistik in gleicher Zeit wie mit dem Koinzidenztrigger erreicht.

8.1.3. GTU-Trigger und Koinzidenz-Trigger kombiniert

Die beiden vorgestellten Trigger bieten jeweils Vor- und Nachteile. Während der GTU-Trigger eine höhere Triggerrate hat, können mit dem Koinzidenztrigger auch die Spuren von Teilchen aufgenommen werden, die zwei oder mehr Stapel durchqueren. Diese Spuren sind wichtig, um die Verschiebungen der Stapel zueinander berechnen zu können.

Mit dem neuen Triggersystem ist es möglich, beide Trigger gleichzeitig zu benutzen, da die L1ctb nicht direkt als Level 1 an das Supermodul weitergegeben wird. Dadurch bleibt die endgültige Entscheidung, ob ein Level 1 ausgegeben wird oder nicht, bei der Triggereinheit. So kann der Triggerbeitrag der GTU bei allen Sequenzen, die durch eine Koinzidenz gestartet wurden, überschrieben werden. Um dazu in der Lage zu sein, müssen die Sequenzen innerhalb der Triggereinheit voneinander unterschieden werden.

Um einen Pretrigger durch Koinzidenz zu erzeugen, wird mehr Logik und somit mehr Zeit benötigt. Außerdem beinhaltet eine Koinzidenz auch immer ein Signal von Bottom. Damit nicht alle Pretrigger durch Signale von Bottom erzeugt werden, werden diese um zwei Takte verzögert. Dadurch werden beide Bedingungen gleichzeitig überprüft. Das hat zwei Vorteile:

1. Innerhalb der Triggereinheit kann zwischen Koinzidenz und GTU-Trigger unterschieden werden.
2. Der Pretrigger wird immer zum gleichen Zeitpunkt relativ zum Szintillatorsignal des eigentlichen Teilchens erzeugt. Dadurch ist die Signalform für beide Trigger gleich und die Rekonstruktion der Daten wird erheblich vereinfacht.

Wird ein Pretrigger durch Koinzidenz erzeugt, wird ein Bit gesetzt, das die nächste Level 1-Entscheidung überschreibt. Dadurch wird für alle Triggersequenzen, die durch Koinzidenz gestartet wurden, auch dann ein Level 1 ausgegeben, wenn die GTU keine Spur in einem Stapel gefunden hat.

Indem beide Trigger gleichzeitig betrieben werden, können die Spuren von Teilchen, die mehr als einen Stapel durchqueren, aufgenommen werden und gleichzeitig eine hohe Statistik durch den GTU-Trigger erreicht werden.

Bei dem kombinierten Trigger können auch einzelne Szintillatorgruppen ausmaskiert werden.

Performance des kombinierten Triggers

Um die Performance des kombinierten Triggers zu überprüfen, stand das Supermodul XI zur Verfügung. In diesem werden keine Module in Stapel 2 verbaut¹, da dahinter PHOS installiert wird. Der Stapel wird ausgelassen, damit die Strahlungslänge vor PHOS minimiert wird. Supermodul XI besteht somit nur aus vier Stapeln. Die gemessenen Triggerraten können also nur bedingt mit den anderen Raten, die mit einem vollständigen Supermodul bestimmt wurden, verglichen werden. Top wurde für die Messung zwischen Stapel 0 und Stapel 1 platziert.

Die Pretriggerrate liegt ähnlich wie beim GTU-Trigger bei ca. 1030 Signalen pro Sekunde und entspricht der Level 0-Rate. Die Level 1-Rate liegt bei ca. 217 Signalen pro Sekunde. Um zu differenzieren, wieviele von diesen Ereignissen auf Koinzidenzen und wieviele auf den GTU-Trigger zurückgehen, wurden die Anzahl der Level 1-Trigger, die Anzahl der überschriebenen Level 1 und die Anzahl der überschriebenen Level 1, für die die GTU keine L1ctb ausgegeben hat, für eine Stunde gezählt. Die Ergebnisse sind in Tabelle 8.1 zusammengefasst.

Daraus gehen verschiedene Beobachtungen hervor. Der Anteil der Trigger, die ausschließlich durch den GTU-Trigger detektiert werden, liegt bei 67,7 %. Dieser hohe Anteil lässt sich darauf zurückführen, dass durch den Koinzidenztrigger hauptsächlich Trigger für den Stapel generiert werden,

¹ Die fünf Stapel in einem Supermodul werden mit Stapel 0 bis Stapel 4 bezeichnet.

Quelle	Zählrate (/3600 s)	Anteil (%)
Level 1 total	781907	100
Level 1 überschrieben	252527	32,3
Level 1 überschrieben und keine L1ctb	70344	9,0

Tab. 8.1.: Die Anzahl der generierten Level 1. In der ersten Zeile ist die Summe aller Level 1 zu sehen. In der zweiten Zeile sind die Level 1 zu sehen, bei denen die Koinzidenzbedingung erfüllt war und in der dritten Zeile sind die Level 1 zu sehen, die ausschließlich auf Grund der Koinzidenzbedingung ausgegeben wurden. Die Anteile der dritten Spalte beziehen sich auf die Zahl aller Level 1. Es stellt sich heraus, dass 23,3 % der Level 1 durch beide Trigger erzeugt wurden. 9,0 % der Level 1 wurden ausschließlich durch den Koinzidenztrigger generiert, während 67,7 % ausschließlich durch den GTU Trigger generiert wurden.

über dem Top platziert ist, während der GTU-Trigger für das gesamte Supermodul gleichmäßig Trigger generiert.

23,3 % aller Level 1 werden von beiden Triggern generiert. Der Anteil der ausschließlich durch den Koinzidenztrigger generiert wird, liegt bei 9,0 %.

Mit den gemessenen Raten können die einzelnen Raten des Koinzidenz- und GTU-Triggers berechnet werden. Sie sind in Tabelle 8.2 zusammengefasst. Die Rate des GTU-Triggers liegt im Bereich der Erwartung für vier Stapel. Da die Module in Stapel 2 kleiner sind und 25 % weniger Auslesekanäle haben als in den anderen Stapeln, wird erwartet, dass die gemessene Level 1-Rate in der Summe um ca. 16 % geringer ist als bei einem vollständigen Supermodul. Die gemessene Rate ist mit 198 Signalen pro Sekunde um 32 Signale bzw. 14 % geringer (vgl. Kap. 8.1.2), was gut mit der Erwartung übereinstimmt.

Die Rate des Koinzidenz-Triggers ist mit 70 Signalen pro Sekunde jedoch bedeutend geringer als die zuvor in Kap. 8.1.1 bestimmte Rate von 105 Signalen pro Sekunde. Die Koinzidenzrate sollte nicht durch den fehlenden Stapel beeinflusst werden, da sie unabhängig vom Supermodul ist. Die Ursache, dass sie so viel tiefer ist, liegt hauptsächlich darin, dass Top bei der Messung anders positioniert war. Da die Winkelverteilung der kosmischen Strahlung proportional zu $\cos^2 \theta$ [Ams08] ist, erfüllen mehr kosmische Teilchen die Koinzidenzbedingung, wenn Top mittig über dem Supermodul platziert ist. Dann wird der maximale symmetrische Winkelbereich um den steilsten Winkel abgedeckt.

Außerdem ist zu beachten, dass bei dieser Messung zusätzlich zur Überprüfung des GTU-busy noch 500 μs gewartet wurde², um vom Status *Readout* in *Idle* zu wechseln. Dadurch und durch die allgemein höhere Level 1-Rate beim kombinierten Trigger erhöht sich die Totzeit im Vergleich zur Abschätzung der Koinzidenz-Triggerrate aus Kap. 8.1.1.

² Diese Einstellung wurde nur zum Testen des neuen Triggers gewählt. Mittlerweile ist diese Funktion standardmäßig deaktiviert.

Quelle	Zählrate (/s)	Anteil an allen Level 1 (%)
GTU	198	91,0
Koinzidenz	70	32,3
GTU aber nicht Koinzidenz	147	67,7
Koinzidenz aber nicht GTU	20	9,0
Koinzidenz und GTU	51	23,3
Gesamt	217	100

Tab. 8.2.: Eine Auflistung der einzelnen Raten, die berechnet wurden. Auffällig ist die niedrige Koinzidenzrate im Vergleich zur vorher bestimmten Rate. Die Rate ist hier geringer, da Top anders positioniert war.

Abschätzung der Level 1–Raten für ein vollständiges Supermodul

Die vorliegenden Raten sind nur gültig für Supermodule, in denen Stapel 2 fehlt. Im Folgenden soll die Rate für ein vollständiges Supermodul abgeschätzt werden.

Wie bereits diskutiert, kann die Koinzidenzrate erhöht werden, indem Top über Stapel 2 positioniert wird. Für ein vollständiges Supermodul ist diese Position sinnvoll, da so die Rate maximiert wird. Ist Top über Stapel 2 platziert, wird für den Koinzidenztrigger eine Triggerrate von ca. 105 Signalen pro Sekunde erwartet (vgl. Kap. 8.1.1).

Die Rate des GTU–Triggers steigt durch einen weiteren Stapel. Allerdings darf hier nicht linear mit den Stapeln extrapoliert werden, da Stapel 2 kleiner ist als die anderen Stapel. Wenn der GTU–Trigger einzeln betrieben wird, liegt die Rate bei ca. 230 Signalen pro Sekunde. Diese wird auch als Grundlage für den kombinierten Betrieb angenommen.

Das ergibt eine Gesamtrate von ca. 335 Signalen pro Sekunde. Der Anteil, der durch beide Trigger generiert wird, liegt bei der Testmessung bei ca. 23,3 %. Da Top bei einem vollständigen Supermodul mittig über diesem platziert ist, erhöht sich der Bereich, der durch Koinzidenz– und GTU–Trigger abgedeckt wird, leicht. Dadurch steigt auch der Anteil der Level 1, die von beiden Triggern generiert werden.

Nimmt man einen gemeinsamen Anteil von 25 % an, liegt die erwartete Rate im kombinierten Betrieb bei mehr als 250 Signalen pro Sekunde. Das entspricht einer Erhöhung um ca. 20 Signale pro Sekunde im Vergleich zum GTU–Trigger. Außerdem werden durch den Koinzidenztrigger auch Spuren von kosmischen Teilchen aufgenommen, die durch mehr als einen Stapel fliegen.

8.2. Testtrigger

Während der Montage der Supermodule werden verschiedene Tests durchgeführt, um die Funktion der eingebauten Module zu überprüfen. Einige dieser Tests benötigen dazu Trigger, die durch die neue Triggereinheit zur Verfügung gestellt werden.

- *Noise*: Hierbei handelt es sich um einen Test, bei dem das Rauschen der einzelnen Auslese-

kanäle des Supermoduls aufgenommen wird.

- *Stresstest*: Das ist ein Test, bei dem die Datennahme mit hohen Triggerraten simuliert wird.
- *Single Trigger*: Hiermit können einzelne Triggersequenzen erzeugt werden.

8.2.1. Noise

Diese Testmessung wird mit jeder Lage durchgeführt, um sicherzustellen, dass das Rauschen der einzelnen Kanäle der MCMs nicht zu stark ist. Dazu werden die einzelnen Kanäle ausgelesen, obwohl keine Teilchen das Supermodul passiert haben.

Um diesen Test durchzuführen, ist es wichtig, dass der *zero suppression*-Filter in den MCMs deaktiviert ist, da er gerade das Rauschen unterdrückt. Um das Rauschen aufzunehmen, werden Triggersequenzen mit 1000 Hz erzeugt. Dabei werden nach einem Pretrigger automatisch Level 0 und Level 1 erzeugt.

Um die Wahrscheinlichkeit zu verringern, dass ein kosmisches Teilchen das Supermodul passiert, wenn das Rauschen aufgenommen wird, wird ein Veto mit den Szintillatoren generiert. Ist der Noise-Trigger aktiviert und einer der Szintillatoren liefert ein Signal, werden alle Pretrigger für die nächsten $3,5\ \mu\text{s}$ unterdrückt. Da die Driftzeit der Elektronen in den Kammern etwa $2\ \mu\text{s}$ beträgt, wird das detektierte kosmische Teilchen von der Messung ausgeschlossen.

Außerdem werden laufende Triggersequenzen abgebrochen, wenn einer der Szintillatoren ein Signal liefert und die Zustandsmaschine wechselt unmittelbar in *Readout*. Hier wartet die Trigger-einheit, bis das GTU-busy auf low geht und wechselt dann über *Clear* in *Idle*. Das $3,5\ \mu\text{s}$ lange Veto wird unabhängig vom Status der FSM aktiviert, so dass neue Pretrigger verhindert werden, bis mindestens diese Vetozeit abgelaufen ist. Die Gefahr, dass nach Abbruch einer Sequenz und vor Ablauf des Vetos ein neuer Pretrigger ausgegeben werden soll, besteht im Normalfall aber gar nicht, da bei der periodischen Rate von 1 kHz maximal jede Millisekunde ein Pretrigger erzeugt werden kann.

Um das Rauschen zu beurteilen, reichen wenige Ereignisse (< 50) aus. Dadurch, dass nur so wenige Ereignisse aufgenommen werden, ist die Wahrscheinlichkeit, dass zufällig ein kosmisches Teilchen das Supermodul im richtigen Moment passiert, relativ gering. Durch die oben genannten Maßnahmen wird diese Wahrscheinlichkeit noch weiter minimiert.

Eigentlich liegt die Triggerrate für PT, L0 und L1 bei einem 1 kHz. Indem aber Pretrigger unterdrückt und Sequenzen abgebrochen werden, reduzieren sich diese Raten. Um die Raten zu bestimmen, wurden PT, L0 und L1 für zehn Minuten gezählt. Die Ergebnisse sind in Tab. 8.3 zusammengefasst.

Die Pretriggerate liegt bei 991 Signalen pro Sekunde, die L0-Rate bei 988 Signalen pro Sekunde und die L1-Rate bei 972 Signalen pro Sekunde. Bei dieser Messung muss berücksichtigt werden, dass das GTU-busy in der Triggereinheit überschrieben wurde und die Daten nicht aufgezeichnet wurden, um eine ausreichende Statistik zu erreichen. Da das Rauschen nicht unterdrückt wird,

Trigger	Zählrate (/600 s)	Anteil (%)	Rate (/s)
Pretrigger	594711	99,1	991
Level 0	592991	98,8	988
Level 1	583001	97,2	972

Tab. 8.3.: Die Anzahl der generierten PT, L0 und L1 und der Anteil an der Zählrate (1 kHz), die ohne Veto bzw. Abbruch von Sequenzen erzeugt worden wäre. Es gilt zu beachten, dass für diese Messung das GTU-busy überschrieben wurde und die Daten nicht aufgezeichnet wurden, um eine ausreichende Statistik zu erhalten.

müssen große Datenmengen verarbeitet werden. Dadurch wird die Triggerrate im echten Betrieb nach wenigen Ereignissen, die die Puffer der DAQ vollschreiben, stark eingeschränkt. Sie liegt dann bei wenigen Signalen pro Sekunde.

In Abb. 8.6 wurden die Daten, die für Lage 4 des Supermoduls XI mit dem Noise-Trigger aufgenommen wurden, mit einem Makro von M. Kalisky [Kal11] ausgewertet. Der Trigger wird seit Anfang 2010 erfolgreich eingesetzt.

Fastnoise

Um Noise aufzunehmen, wird die komplette Auslekette des Supermoduls gebraucht. Das heißt, dass alle Fasern für die Optical Readout Boards schon verlegt sein müssen. Falls ein Modul aber getauscht werden muss, müssen die Fasern und alle anderen Versorgungskabel wieder entfernt werden.

Darum wurde von V. Angelov [Ang11] eine spezielle Konfiguration für die TRAPs entwickelt, bei der die Daten des Rauschens über die Netzwerkverbindungen der DCS-Boards ausgelesen werden. Dadurch ist es möglich, das Rauschen auszulesen, ohne die ORI-Fasern zu verlegen. In dieser Konfiguration benötigen die TRAPs nur einen Pretrigger und keinen Level 0 bzw. Level 1.

Dieser Trigger wird erzeugt, indem beim Noise-Trigger die Ausgabe des Level 0 und des Level 1 unterdrückt wird.

8.2.2. Stresstest

In Kap. 8.1 wurden Trigger vorgestellt, mit denen die Spuren kosmischer Teilchen mit dem Supermodul aufgenommen werden können. Die höchsten Triggerraten werden erreicht, wenn der Koinzidenz- mit dem GTU-Trigger kombiniert wird. Allerdings liegt die Level 1-Rate dann immer noch unter 300 Signalen pro Sekunde. Um die Datennahme mit höheren Triggerraten zu simulieren, wurde der Stresstest entwickelt.

Bevor das neue Triggersystem aufgebaut wurde, wurden für den Stresstest einfach periodische Pretrigger mit 1 kHz erzeugt. Die TRAPs wurden so konfiguriert, dass sie keinen Level 0 und keinen Level 1 benötigen.

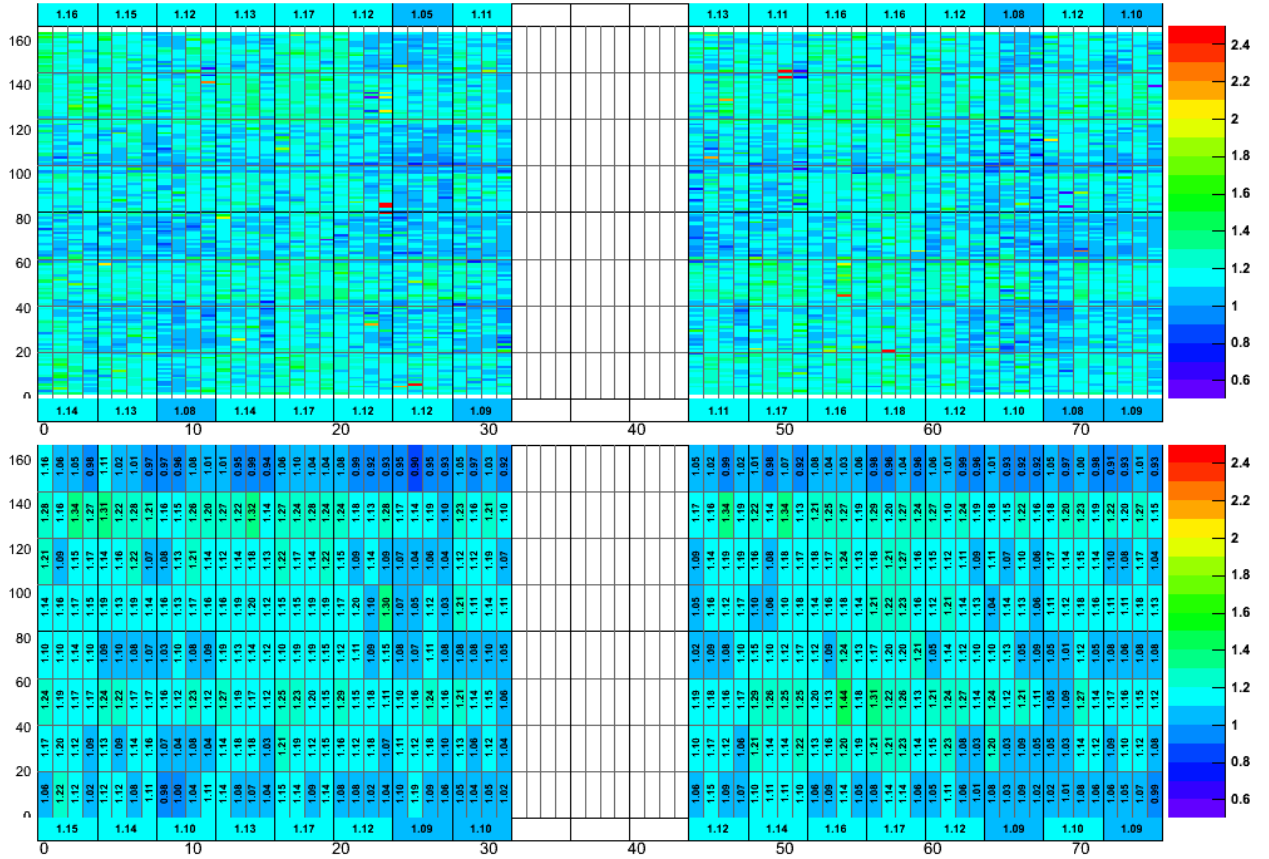


Abb. 8.6.: Darstellung des Rauschens, aufgenommen mit dem Noise-Trigger. Hier ist Lage 4 des Supermoduls XI zu sehen. Das dritte Modul fehlt, da Stapel 2 in diesem Supermodul nicht verbaut wird. Oben ist jeder einzelne Kanal der MCMs als dünner waagerechter Strich zu sehen. Die Stärke des Rauschens kann mit der Farbskala rechts verglichen werden. Die Skala hat eine willkürliche Einheit. Einzelne Kanäle sind rot, was bedeutet, dass sie verhältnismäßig stark rauschen. Bei der großen Anzahl der Kanäle wird das jedoch toleriert. Unten ist das mittlere Rauschen für jeden MCM zu sehen.

Mit der neuen Triggereinheit wurde ein Stresstest entwickelt, der vollständige und zufällige Triggersequenzen erzeugt. Dabei sind die Pretrigger-, Level 0-, und Level 1-Raten frei konfigurierbar. Im Folgenden wird zunächst vorgestellt, wie in einem FPGA (Pseudo-)Zufallszahlen erzeugt werden. Diese werden benötigt, um zufällige Trigger zu erzeugen. Danach wird die Funktionsweise des Stresstests erläutert.

Pseudozufallszahlen

Um zufällige Triggersequenzen zu erzeugen, müssen zunächst Zufalls- bzw. Pseudozufallszahlen erzeugt werden. Dazu wird ein linear rückgekoppeltes Schieberegister (LFSR kurz für *Linear Feedback Shift Register*) verwendet.

In Abb. 8.7 ist ein 11-Bit LFSR und seine Funktionsweise dargestellt. Im Wesentlichen handelt es

sich bei einem LFSR um eine n -Bit breite Zahl, bei der mit jedem Takt jedes Bit um eine Position zum nächst signifikanteren Bit verschoben wird. Für Bit 0 wird dabei das logische XOR, also das exklusive ODER, von ausgewählten anderen Bits gesetzt. Bei geschickter Wahl dieser Eingangs-Bits durchläuft das LFSR alle Zustände und wiederholt sich erst nach $2^n - 1$ Takten, wobei n die Breite des Registers ist.

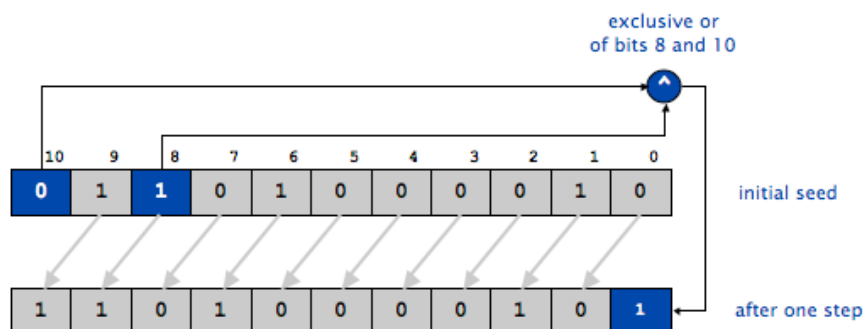


Abb. 8.7.: Ein 11-Bit breites LFSR. Mit jedem Takt werden die Bits um eine Position nach links geschoben. Für Bit 0 wird das logische XOR von Bit 8 und 10 gesetzt. Dadurch werden Pseudozufallszahlen erzeugt. Bild aus [Pri10]

Der Startwert (*initial seed*) darf nicht ausschließlich aus Nullen bestehen, da $0 \text{ XOR } 0 = 0$. Das LFSR würde also in diesem Zustand stehen bleiben. Im Normalfall werden Flip-Flops mit dem Wert Null gestartet. Von daher bietet es sich an, für einen FPGA das negierte exklusive ODER, XNOR, anstelle des XOR zu verwenden, da dann alle Startwerte außer ausschließlich aus Einsen erlaubt sind.

Um die Zufallszahlen in der Triggereinheit zu erzeugen, wird ein 32-Bit breites LFSR verwendet. Als Rückkopplung werden die Bits 31, 21, 1 und 0 mit einem XNOR verknüpft [Xil96]. Dadurch wiederholen sich die Pseudozufallszahlen nach $2^{32} - 1 \approx 4,3 \cdot 10^9$ Takten. Das entspricht ca. 100 s bei einem 40 MHz-Taktgeber, was mehr als ausreichend für den Stresstest ist. Der aktuelle Wert des LFSR wird mit L und der Maximalwert mit $L_m = 2^{32} - 1$ bezeichnet.

Logik des Stresstest

Für den Stresstest wird die normale Triggersequenz wie für andere Trigger durchlaufen. Das Verhalten der Zustandsmaschine wird jedoch ein wenig abgeändert, da alle Signale intern generiert werden. Die Zeitvorgaben für die FSM sind genauso wie bei den anderen Triggern.

In Abb. 8.8 ist das Verhalten der FSM beim Stresstest schematisch dargestellt. Die Bedingungen *Condition PT*, *Condition L0*, *Condition L1* und *Condition RE* definieren, ob der entsprechende Trigger an der jeweiligen Stelle erzeugt wird bzw. wann von *Readout* in *Clear* gewechselt wird. Sie sind in Tabelle 8.4 zusammengefasst. Für jeden Trigger kann eine individuelle Wahrscheinlichkeit eingestellt werden, mit der die Wahrscheinlichkeit festgelegt wird, mit der er an entsprechender

Stelle ausgegeben wird. Sie wird mit $p(\text{Trigger})$ bezeichnet.

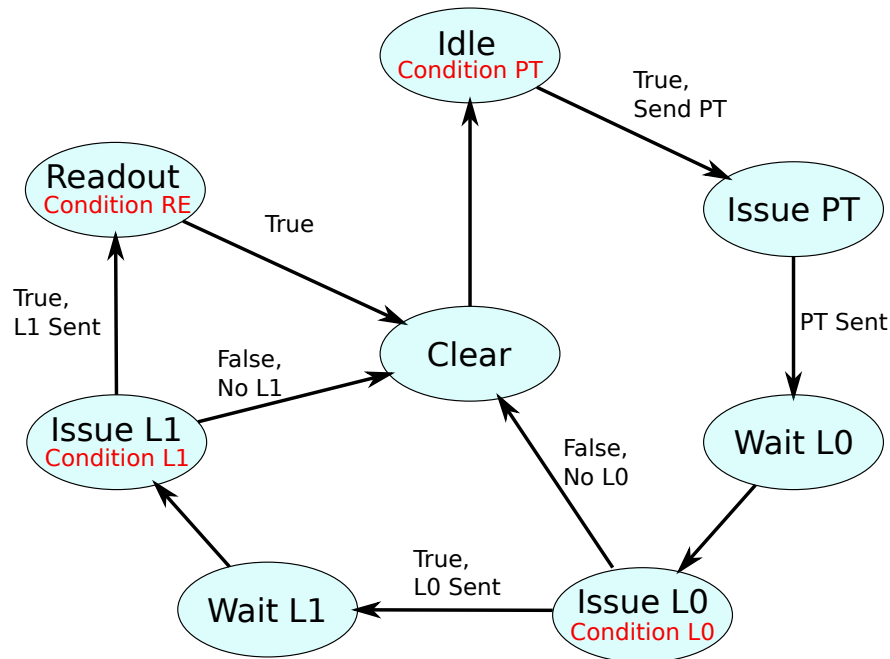


Abb. 8.8.: Verhalten der FSM beim Stresstest. In den entsprechenden Zuständen wird mit Condition PT, L0 und L1 entschieden, ob der jeweilige Trigger ausgegeben oder die Sequenz abgebrochen wird. Mit Condition RE wird festgelegt, wann die FSM von *Readout* in *Clear* wechselt.

	Bedingung	Voreinstellung
Condition PT	$L < p(\text{PT}) \cdot L_m$ Wird mit Frequenz $f = 0,01 \text{ Hz} - 40 \text{ MHz}$ in <i>Idle</i> überprüft	$p(\text{PT}) = 0,5$ $f = 40 \text{ MHz}$
Condition L0	$L < p(\text{L0}) \cdot L_m$ Wird einmalig in Status <i>Issue L0</i> überprüft	$p(\text{PT}) = 0,5$
Condition L1	$L < p(\text{L1}) \cdot L_m$ Wird einmalig in Status <i>Issue L1</i> überprüft	$p(\text{PT}) = 0,5$
Condition RE	Warte Δt oder GTU-busy = 0 Wird solange geprüft bis wahr	$\Delta t = 500 \mu\text{s}$

Tab. 8.4.: Die Bedingungen und Voreinstellungen unter denen ein Trigger ausgegeben, bzw. von *Readout* in *Clear* gewechselt wird. L ist der jeweils aktuelle Wert des LSFR. L_m ist der Maximalwert des LSFR und $p(\text{Trigger})$ entspricht der Wahrscheinlichkeit, mit der der entsprechende Trigger an gegebener Stelle erzeugt wird. Δt entspricht der Zeitvorgabe PT-Clear der FSM. Die tatsächliche Zeit in *Readout* beträgt also ca. $490 \mu\text{s}$.

Um einen Pretrigger zu generieren, werden periodische Pulse mit einer Frequenz zwischen $0,01 \text{ Hz}$ und 40 MHz erzeugt. Damit die Pretrigger zufällig sind, kann eine Wahrscheinlichkeit $p(\text{PT})$ eingestellt werden, mit der ein periodischer Puls tatsächlich als Pretrigger ausgegeben wird. Wird ein Pretrigger ausgegeben, wechselt die FSM automatisch in *Wait L0* und anschließend in *Issue L0*.

In *Issue L0* wird nicht, wie üblich, direkt ein Level 0 generiert, sondern nur mit der eingestellten Wahrscheinlichkeit $p(L0)$. Die Wahrscheinlichkeit entspricht somit dem Verhältnis der L0-Rate zur PT-Rate. Falls kein L0 generiert wird, wechselt die FSM in *Clear*.

Der Level 1 wird analog zum Level 0 generiert. Erreicht die FSM *Issue L1*, wird mit der eingestellten Wahrscheinlichkeit ein Level 1 ausgegeben und die FSM wechselt in *Readout*. Die Wahrscheinlichkeit entspricht somit dem Verhältnis der L1-Rate zur L0-Rate.

In *Readout* gibt es zwei Optionen. Entweder es wird eine feste Zeit ($500 \mu s$) eingestellt oder das GTU-busy wird benutzt um in *Clear* zu wechseln. Da die Daten beim Stresstest in der Regel nicht aufgezeichnet werden, wird die GTU gar nicht benutzt und das GTU-busy ist dauerhaft high. Die erste Variante ist deshalb die Standardeinstellung. Dadurch, dass hier gewartet wird, wird sichergestellt, dass die TRAPs die Rohdaten verschickt haben und bereit für neue Trigger sind.

In Abb. 8.9 ist eine Aufnahme des ILA für den Stresstest zu sehen. Hier wurde auf einen Level 1 getriggert und nur noch 23 neue Takte geschrieben, da nach einem Level 1 für ca. $490 \mu s$ keine neuen Trigger zu sehen sind. Der ILA kann jedoch nur ca. $25 \mu s$ aufnehmen. Indem so getriggert wurde, können die meisten Trigger beobachtet werden.

Es ist zu sehen, wie zunächst ein Pretrigger ohne Level 0 und dann ein Pretrigger mit Level 0 erzeugt wird. Danach wird allerdings kein Level 1 erzeugt. Daraufhin gibt es eine längere Pause, da durch den Level 0 eine Totzeit von ca. $9 \mu s$ entsteht. Im Anschluss ist nach zwei einzelnen Pretriggern eine komplette Sequenz zu sehen.

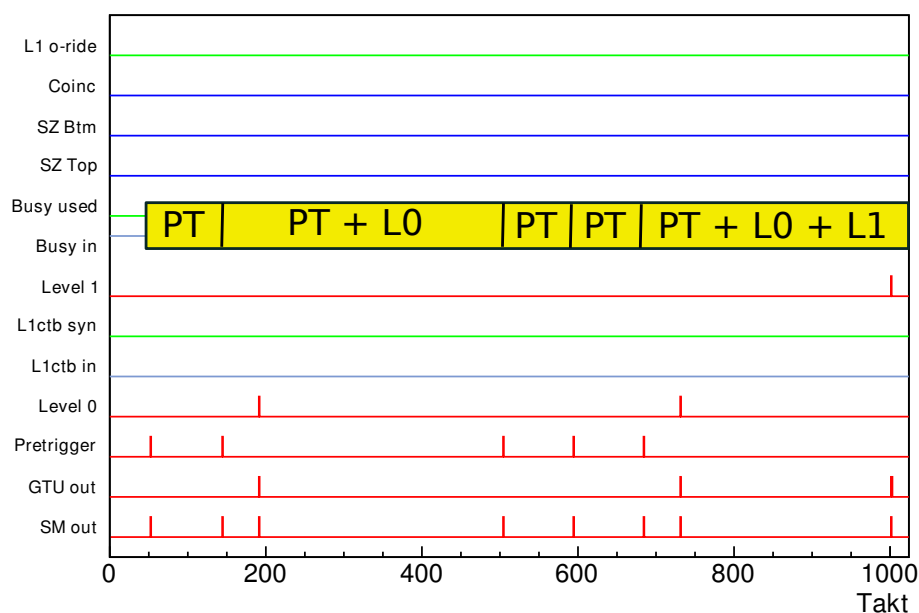


Abb. 8.9.: Hier ist gut zu sehen, dass die Sequenzen nach Pretrigger oder Level 0 abgebrochen werden können. Außerdem sind die Totzeiten nach den entsprechenden Triggern durch die gelben Kästen angedeutet. Die Totzeit nach einem Level 1 ist hier nicht vollständig, da sie ca. $250 \mu s$ lang ist.

Die Triggerraten sind für die jeweiligen Einstellungen der Wahrscheinlichkeiten, der Frequenz und der Pause in *Readout* verschieden. In Tabelle 8.5 sind die Triggerraten für drei verschiedene Einstellungen aufgelistet.

Einstellung	f	40 MHz	40 MHz	40 MHz
	p(PT)	50 %	50 %	50 %
	p(L0)	50 %	50 %	20 %
	p(L1)	50 %	50 %	5 %
	Δt	500 μs	250 μs	500 μs
Triggerraten (10^3 s^{-1})	Pretrigger	7,8	15,1	117,7
	Level 0	3,9	7,6	23,5
	Level 1	1,9	3,7	1,2

Tab. 8.5.: Die Triggerraten in Abhängigkeit der Einstellungen. Sie wurden zur Übersichtlichkeit auf eine Nachkommastelle gerundet. Die Raten des Pretriggers und des Level 0 werden durch eine hohe Level 1-Rate stark reduziert, da die Zeit, die die Triggereinheit nach einem Level 1 in *Readout* verbleibt, vergleichsweise lang ist. Durch die Halbierung von Δt werden die Raten aller Trigger nahezu verdoppelt. Die erste Spalte gibt die Standardeinstellung an. Die eingestellten Verhältnisse zwischen L0 und PT und zwischen L1 und L0 werden durch die Raten bestätigt.

Bei den verschiedenen Einstellungen wird deutlich, dass die Triggerraten am stärksten durch die Einstellung für Δt und dem Verhältnis Level 1 zu Level 0 beeinflusst werden. Das liegt daran, dass in dem Zeitraum Δt nach einem Level 1 keine neuen Sequenzen gestartet werden dürfen.

Die Totzeit bei einer vollständigen Sequenz ist demnach ca. 500 μs bzw. Δt . Wird die Sequenz nach einem Level 0 abgebrochen, liegt die Totzeit bei ca. 9 μs . Wird sie nach einem Pretrigger abgebrochen, liegt die Totzeit nur bei ca. 2 μs .

Δt ist allerdings ein kritischer Wert für die Stabilität des Tests. Wird die Zeit zu kurz gewählt, wird eine neue Triggersequenz gestartet, während die TRAPs noch die Rohdaten ausgeben. Dadurch können die TRAPs zum Absturz gebracht werden.

Da es der Sinn des Stresstests ist, die Supermodule bzw. die Datennahme mit hohen Raten und nicht die Stromversorgung zu testen, muss außerdem darauf geachtet werden, dass die Triggerrate in einem Bereich liegt, in dem die Stromversorgung sicher funktioniert. Deshalb wurde als Standardeinstellung konservativ die erste Spalte gewählt. Dadurch ist ein stabiler Betrieb sichergestellt und alle Probleme, die auftreten, sind tatsächlich auf das Supermodul zurückzuführen. Die gewählten Einstellungen werden seit Mai 2010 verwendet und wurden bisher (Februar 2011) an drei Supermodulen getestet. Die Stromversorgung ist seitdem nicht zusammengebrochen.

Für den eigentlichen Test wurde ein Perl-Skript von T. Dietel [Die11] entwickelt, dass das Supermodul automatisch konfiguriert, den Trigger jeweils für 10 Minuten anstellt und in den Pausen alle TRAPs kontrolliert. Eventuelle Fehler werden in einer Datei protokolliert. Dadurch kann der Stresstest über lange Zeiten betrieben werden, ohne dass er überwacht werden muss.

Falls es gewünscht ist, kann das Skript erweitert werden, so dass verschiedene Konfigurationen bzw. Raten während eines Stresstests durchgeführt werden. Dazu müssen lediglich die Befehle zum

Ändern der Raten in das Skript aufgenommen werden (vgl. Kap. 9.2.3).

8.2.3. Einzelne Sequenzen

Zu Testzwecken ist es manchmal nötig, nur eine Sequenz eines Triggers auszugeben. Dieser Modus kann für alle Trigger mit einem zusätzlichen Bit bei der Triggerauswahl gewählt werden.

Ist es aktiviert, werden zunächst keine Trigger ausgegeben. Erst wenn von außen ein Start-Bit gesetzt wird, wird der ausgewählte Trigger gestartet. Sobald die erste Triggersequenz beendet wird, also ein Level 1 ausgegeben wird, wird der Trigger wieder gestoppt. Das bedeutet auch, dass unter Umständen mehrere Pretrigger und Level 0 ausgegeben werden.

Der Modus ist insbesondere sinnvoll für Trigger, bei denen tatsächlich Daten aufgenommen werden. Das sind der Koinzidenz-, GTU- und Noise-Trigger. Indem nur ein Ereignis gespeichert wird und der Trigger anschließend gestoppt wird, können die Daten z.B. direkt in der GTU eingesehen werden, ohne eine Rekonstruktion der Daten vorzunehmen.

Einzelne Pretrigger

Zusätzlich zu den Triggern, die mit der FSM erzeugt werden, können noch einzelne Pretrigger-Sequenzen erzeugt werden.

Über Register kann der Trigger konfiguriert werden. Dabei kann eingestellt werden, wieviele Trigger und in welchem Abstand sie ausgegeben werden sollen. Die Pulse haben eine konstante Länge von 25 ns und den gleichen eingestellten Abstand zueinander.

Um die Pretrigger-Sequenz zu starten, wird das gleiche Start-Bit benutzt wie bei den anderen einzelnen Sequenzen. Dieser Trigger wurde für einige sehr spezielle Tests verwandt und wird nur der Vollständigkeit halber erwähnt.

In diesem Kapitel wurden die verfügbaren Trigger vorgestellt. Dazu zählen der Koinzidenz- und der GTU-Trigger, mit denen auf kosmische Teilchen getriggert werden kann. Der GTU-Trigger hat dabei eine höhere Triggerrate, dafür wird aber in der Regel nicht auf Teilchen getriggert, die durch zwei Stapel des Supermoduls fliegen. Indem die beiden Trigger kombiniert wurden, werden die Vorteile von beiden ausgenutzt. Insgesamt liegt die Zählrate mit ca. 250 Signalen pro Sekunde höher als die des GTU-Triggers, während durch den Koinzidenztrigger Spuren nachgewiesen werden, die durch mehr als einen Stapel gehen.

Zum Testen des Supermoduls ist ein konfigurierbarer Stresstest-Trigger entwickelt worden, der zufällige Triggersequenzen erzeugt. Außerdem kann das Rauschen der einzelnen Kanäle des Supermoduls mit dem Noise- bzw. Fastnoise-Trigger aufgenommen werden.

Allgemein bleibt zu vermerken, dass alle Trigger gut funktionieren und es keine Probleme mit

der Stromversorgung oder abstürzenden TRAPs mehr gibt, seitdem die Zustandsmaschine in der Triggereinheit zur Triggergeneration genutzt wird.

Die Spuren kosmischer Teilchen, die mit dem neuen Triggersystem aufgenommen wurden, wurden in [Alb10], [Gat10] und [Wal10] benutzt. Die Testtrigger werden für jedes Supermodul gebraucht.

9. Benutzerschnittstelle

In diesem Kapitel wird die Benutzerschnittstelle vorgestellt, mit der alle Funktionen der Triggereinheit überwacht und kontrolliert werden können. Zunächst wird erläutert, wie mit dem FPGA kommuniziert werden kann. Danach werden die einzelnen Komponenten der Benutzerschnittstelle vorgestellt. Im Anschluss werden die verfügbaren Befehle, mit denen die Triggereinheit kontrolliert wird, diskutiert.

9.1. Kommunikation mit dem FPGA

Zugriffe auf den FPGA der Triggereinheit erfolgen über die VME-Schnittstelle. In das Design für den FPGA wurden 16-Bit breite Register eingebaut, die über den VME-Bus gelesen und/oder geschrieben werden können.

Jedem Register wird eine hexadezimale Adresse zugewiesen. Sie setzt sich zusammen aus der gemeinsamen vierstelligen Adresse des V1495 0x3210 und einer eindeutigen vierstelligen Adresse im Bereich von 0x0000 bis 0x7FFC.

Für jedes Register ist festgelegt, ob es gelesen und/oder geschrieben werden kann. Im FPGA können alle Register gelesen und verarbeitet werden. Auf die Register, die nicht über VME geschrieben werden können, hat der FPGA außerdem Schreibzugriff.

Die insgesamt 80 Register sind in Anhang E mit den entsprechenden Adressen und einer Kurzbeschreibung aufgelistet.

9.1.1. Motivation für eine Benutzerschnittstelle

Über die VME-Schnittstelle ist das V1495 mit dem Computer DAQ00 verbunden. Mit den Kommandozeilenprogrammen *vme_read_d16* und *vme_write_d16* können die 16-Bit breiten Register gelesen bzw. geschrieben werden. Dafür muss jedoch bei jedem einzelnen Zugriff die Adresse des jeweiligen Registers bekannt sein und eingegeben werden. Erschwerend sind im Design einige 32-Bit breite Register eingebaut, für die zwei Zugriffe erforderlich sind.

Da der Trigger von vielen verschiedenen Personen konfiguriert und überwacht werden soll, soll der Zugriff so einfach wie möglich erfolgen, damit keine lange Einarbeitungszeit nötig ist. Durch eine geeignete Schnittstelle wird die Konfiguration außerdem erheblich beschleunigt.

9.2. Aufbau der Benutzerschnittstelle

In Abb. 9.1 ist der Aufbau für die Kommunikation mit der Triggereinheit skizziert. Die Triggereinheit ist über den VME-Bus mit dem Computer DAQ00 verbunden. Dort sind zwei Programme vorhanden, die auf den FPGA zugreifen können. In 7.4.2 wurde bereits das Kontrollprogramm für den ILA vorgestellt. Zusätzlich wurde ein Server programmiert, der als Zwischenebene zwischen Benutzer und Triggereinheit fungiert. Dadurch ist es möglich, die Triggereinheit über ein TCP/IP-Netzwerk mit einem Client zu überwachen und zu kontrollieren.

Um die Prozesse weiter zu vereinfachen, werden einige Kontrolldaten im *SuperModule Monitor* (SMMon) auf dem Computer TRD01 angezeigt, mit dem auch alle anderen Systeme in der Montagehalle überwacht werden. Außerdem wurde das Skript *trigger* geschrieben, das auf den Server und das ILA-Steuerprogramm zugreift, um alle Funktionen der Triggereinheit in einem Befehl zu bündeln. Die einzelnen Komponenten werden im Folgenden vorgestellt.

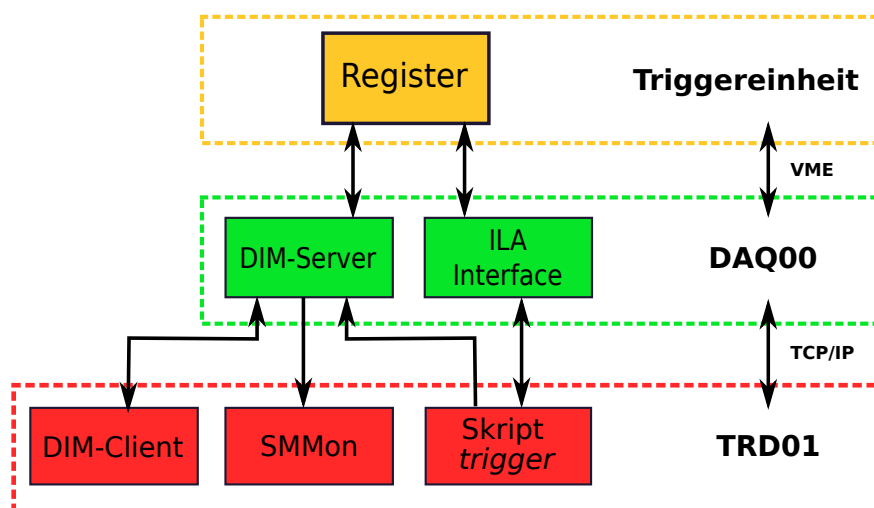


Abb. 9.1.: Hier ist die Benutzerschnittstelle skizziert. Die DAQ00 ist über die VME-Schnittstelle mit der Triggereinheit verbunden. Auf der DAQ00 liegt das Kontrollprogramm für den ILA. Außerdem wurde ein DIM-Server programmiert, der die Register in der Triggereinheit schreiben und auslesen kann. Über ein TCP/IP-Netzwerk nimmt er Befehle entgegen und veröffentlicht ausgelesene Register. Der Server kann von jedem Rechner im Netzwerk angesprochen werden. In der Montagehalle wird dazu in der Regel der Computer TRD01 verwendet. Auf diesem wird dazu ein DIM-Client benutzt. Zusätzlich werden noch einige Daten im SMMon angezeigt. Mit dem Skript *trigger* können alle Funktionen der Triggereinheit eingestellt werden und der ILA kontrolliert werden.

9.2.1. DIM-Server

Als Zwischenebene zwischen dem Benutzer und der Triggereinheit, wurde ein DIM-Server programmiert. Das *Distributed Information Management System* (DIM) ist ein Kommunikationssystem, das auf dem Client/Server-Prinzip aufbaut und am CERN entwickelt wird.

In Abb. 9.2 ist eine schematische Übersicht über die Funktionsweise von DIM dargestellt. Alle Server registrieren sich bei einem *Name Server* mit ihren *Services*. Ein Client kann diese Informationen der verfügbaren Server abrufen und kommuniziert dann direkt mit dem gewünschten Server.

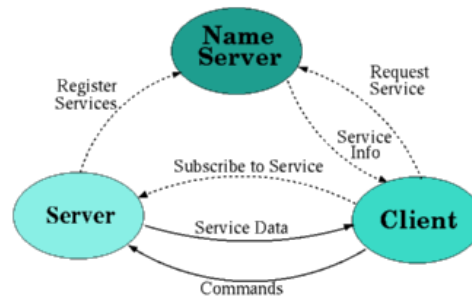


Abb. 9.2.: Funktionsprinzip von DIM. Jeder Server registriert sich bei einem Name Server, der dann alle Services des Servers auflistet. Ein Client kann auf diese Liste zugreifen und kommuniziert dann direkt mit dem entsprechenden Server. Außerdem können Befehle über einen Command Channel an den Server geschickt werden. Bild aus [DIM11]

Ein Server veröffentlicht verschiedene Services, die abonniert werden können. Sobald ein Service vom Server aktualisiert wird, wird er beim Client, der den Service abonniert hat, auch aktualisiert. Dadurch können Kontrollparameter einfach überwacht werden. Außerdem können *Command Channels* veröffentlicht werden, die Befehle entgegennehmen. Für weitere Informationen über DIM sei auf die Homepage des Projekts [DIM11] verwiesen.

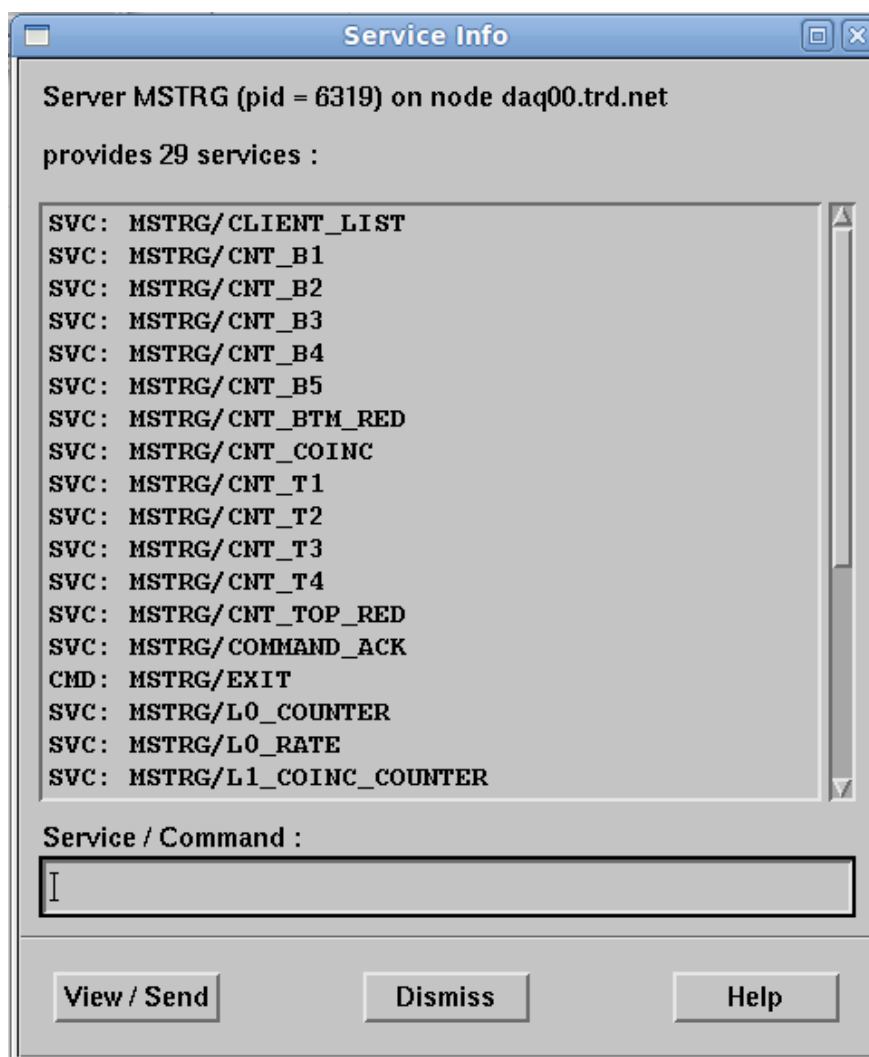
Der DIM-Server, der für die Triggereinheit programmiert wurde, registriert sich mit dem Namen *mstrg* beim Name Server. In Abb. 9.3 ist eine Abbildung eines DIM-Clients zu sehen, der mit dem Server verbunden ist. Der Server liest jede Sekunde Parameter der Triggereinheit aus. Dazu zählen die Zähler für die einzelnen Szintillatoren, Trigger und Koinzidenz, die aktuellen Triggerraten und der gewählte Trigger. Falls in Zukunft weitere Register der Triggereinheit überwacht werden sollen, lassen sich die entsprechenden Services einfach in den Server einbauen.

Um den Verkehr im Netzwerk zu minimieren, werden die entsprechenden Services nur dann aktualisiert, wenn sich der Wert tatsächlich geändert hat.

Über den Command Channel *mstrg/command* werden Befehle, um die Triggereinheit zu steuern, entgegengenommen. Der DIM-Server verarbeitet die eingehenden Befehle und schreibt die entsprechenden Werte in die entsprechenden Register. Die einzelnen Befehle werden mit dem Skript *trigger* in Kap. 9.2.3 diskutiert, da der DIM-Server nur eine Zwischenebene darstellt.

9.2.2. Überwachung der Triggereinheit

Um den Trigger während des laufenden Betriebs zu überwachen, wurde ein Modul für den *Supermodule Monitor* entwickelt, das einige Informationen der Triggereinheit laufend darstellt.

Abb. 9.3.: Screenshot eines DIM-Clients der auf *mstrg* zugreift.

Im SMMon werden Daten des Supermoduls und von den Versorgungssystemen gesammelt und gebündelt dargestellt. Dadurch lassen sich alle Systeme mit einem Programm überwachen. Der SMMon greift dazu auf die verschiedenen DIM-Server zu und abonniert die entsprechenden Services.

In Abb. 9.4 ist ein Screenshot des SMMon-Moduls zu sehen, das die Informationen der Triggereinheit darstellt. Da im SMMon relativ viele Informationen in einem Fenster dargestellt werden, ist der verfügbare Platz beschränkt. Aus diesem Grund werden nur einige Services der Triggereinheit abonniert. Das sind im Einzelnen die Triggerraten (*Rate*) und die absoluten Zahlen (*CNT*) von Pre-trigger, Level 0 und Level 1. Die Zähler lassen sich über einen Befehl zurücksetzen. Da die Zähler der Überwachung dienen, können sie nicht gestoppt werden. Sobald die Triggereinheit eingeschaltet ist, laufen auch die Zähler. Am rechten Rand wird der ausgewählte Trigger (*Mode*) angezeigt. Dazu wird einmal der Wert des Auswahlregisters dargestellt und für bekannte Werte des Registers

außerdem ein Kürzel für den jeweiligen Trigger angezeigt (vgl. Kap. 9.2.3).

Alle anderen Kontrollparameter der Triggereinheit können mit einem DIM-Client beobachtet werden. Falls es in Zukunft nötig sein sollte, sie in den SMMon einzubinden, kann das Modul relativ einfach um weitere Services ergänzt werden.

RATE	PT	7960/s		L0	3961/s		L1	1921/s		MODE	#	21
CNT		13730197			6867064			3425403				stress

Abb. 9.4.: Screenshot des SMMon-Moduls. In der oberen Zeile werden die aktuellen Raten der Trigger angezeigt. In der unteren sind die absoluten Zähler der einzelnen Trigger zu sehen. Rechts wird der ausgewählte Trigger angezeigt. Hier ist der Stresstest aktiviert.

9.2.3. Kontrollskript

Das Skript *trigger*, das auf der TRD01 mit dem Befehl *trigger* aufgerufen wird, fasst alle Konfigurationsmöglichkeiten der Triggereinheit und die Bedienung des ILA zusammen. Dazu greift es auf den Befehlskanal *mstrg/command* des DIM-Servers zu. Die einzelnen Optionen sind in Tabelle 9.1 zusammengefasst und werden in den folgenden Unterkapiteln vorgestellt. Für die einzelnen Optionen wurden außerdem Hilfen geschrieben, die mit *trigger help <Option>* aufgerufen werden können.

Option	Beschreibung
s	Triggerauswahl
t	Konfiguration der Zeitvorgaben der FSM
reset	Zurücksetzen der Triggerzähler
stress	Konfiguration des Stresstests
mask	Ausmaskieren von Szintillatorgruppen
ila	Alle Optionen des ILA
sng_pt	Konfiguration des Pretrigger-Sequenz-Triggers
sng_shot	Startet einzelne Triggersequenzen
sz_cnt	Zum Starten, Stoppen und Zurücksetzen der Szintillatorzähler
info	Schreibt komplette Konfiguration in Datei
help	Hilfe anzeigen

Tab. 9.1.: Die Optionen, die beim Skript *trigger* zur Verfügung stehen. Sie werden ausgeführt mit *trigger <Option>*. Eine Hilfe zu den einzelnen Funktionen erhält man über *trigger help <Option>*.

Triggerauswahl

Um einen Trigger auszuwählen, gibt es zwei verschiedene Methoden. Das Skript *trigger* nimmt einen Zahlenwert für das entsprechende Register entgegen. Außerdem akzeptiert es die Kürzel von bekannten Triggern. Falls neue Kürzel eingeführt werden sollen, können sie im Skript eingefügt werden. Die verfügbaren Trigger sind in Tabelle 9.2 aufgelistet. Sie wurden bereits in Kap. 8 diskutiert.

Wert	Kürzel	Beschreibung
0	coinc	Koinzidenz-Trigger
1	noise	Noise-Trigger
2	btm_gtu	GTU-Trigger
3	off	Kein Trigger / Deaktiviert
4	sng_pt_seq	Pretrigger-Sequenz
5	stress_gtu_busy	Stresstest-Trigger mit GTU-Busy
9	coinc_btmgtu	Koinzidenz- und GTU-Trigger kombiniert
21	stress	Stresstest-Trigger
48	coinc_sng	Für eine einzelne Koinzidenz-Trigger Sequenz. Ignoriert GTU-busy
49	noise_sng	Für eine einzelne Noise-Trigger Sequenz. Ignoriert GTU-busy
50	btmgtu_sng	Für eine einzelne GTU-Trigger Sequenz. Ignoriert GTU-busy
214	fnoise	Fastnoise-Trigger

Tab. 9.2.: Die verfügbaren Trigger. Über die zusätzlichen Bits können die Trigger weiter konfiguriert werden.

Für die eigentliche Triggerauswahl stehen die ersten vier Bits des Registers zur Verfügung. Die anderen Bits aktivieren Zusatzeinstellungen, die unabhängig vom gewählten Trigger sind. Deren Funktionen sind in Tabelle 9.3 zusammengefasst. Um sie zu aktivieren, muss ihr Wert zum jeweiligen Trigger addiert werden. Für einige der oben genannten Trigger sind einige Bits standardmäßig aktiviert.

Bit	Wert	Beschreibung
4	16	Das GTU-busy wird ignoriert
5	32	Einzelne Triggersequenzen
6	64	Level 0 der Triggersequenzen werden nicht ausgegeben
7	128	Level 1 der Triggersequenzen werden nicht ausgegeben
8	256	Überschreibt alle Level 1-Entscheidungen intern
9	512	Aktiviert die Verzögerung im Status <i>Readout</i>

Tab. 9.3.: Die Werte und Funktionen der Zusatzbits zur Triggerauswahl. Sie können unabhängig vom gewählten Trigger aktiviert werden. Für einige Trigger aus Tab. 9.2 sind einige Bits standardmäßig aktiviert.

Zeitvorgaben

Die Zeitvorgaben der Zustandsmaschine sind einstellbar. Allerdings muss beachtet werden, dass sie auf die Konfigurationen der TRAPs und der GTU abgestimmt sind. Bei einer Änderung funktionieren die Trigger unter Umständen nicht mehr. Um die Zeitvorgaben zu ändern, wird *trigger t* *<Option>* *<Wert>* benutzt. In Tabelle 9.4 sind die einzelnen Optionen gelistet.

Option	Beschreibung
l0	Zeit zwischen Pretrigger und Level 0, in TTC-Takten
l1	Zeit zwischen Pretrigger und Level 1, in TTC-Takten
read	Verzögerung im Status <i>Readout</i> der FSM, in TTC-Takten
clear	Verzögerung im Status <i>Clear</i> der FSM, in TTC-Takten
reset	Setzt alle Zeiteinstellungen auf Standard zurück

Tab. 9.4.: Die Optionen, um die Zeitvorgaben der Zustandsmaschine anzupassen.

Stresstest-Trigger Konfiguration

Um den Stresstest bzw. die einzelnen Triggerraten zu konfigurieren, wird der Befehl *trigger stress* *<Option>* *<Wert>* benutzt. Die einzelnen Optionen sind in Tabelle 9.5 aufgelistet.

Option	Beschreibung
freq	Häufigkeit, mit der potentielle PT erzeugt werden Der Eingabewert entspricht der Anzahl von TTC-Takten zwischen zwei Pulsen
ptfreq	Wahrscheinlichkeit (in Prozent), mit der ein potentieller PT erzeugt wird
l0pt	Wahrscheinlichkeit (in Prozent), mit der ein L0 erzeugt wird Entspricht dem Verhältnis der Level 0-Rate zur Pretrigger-Rate
l1l0	Wahrscheinlichkeit (in Prozent), mit der ein L1 erzeugt wird Entspricht dem Verhältnis der Level 1-Rate zur Level 0-Rate

Tab. 9.5.: Die Optionen, um den Stresstest-Trigger zu konfigurieren

Maske für die Szintillatorgruppen

Die einzelnen Szintillatorgruppen werden bit-weise mit einem logischen UND einer Maske verknüpft. Diese besteht dementsprechend aus einem 9 Bit breiten Vektor. Standardmäßig sind alle Szintillatoren aktiviert. Das bedeutet, dass alle neun Bits auf high stehen. In Tabelle 9.6 sind die Gruppen den einzelnen Bits zugeordnet. Um nur bestimmte Gruppen zu aktivieren, müssen die entsprechenden Werte aufsummiert werden. Die Summe wird dann mit *trigger mask* *<Summe>* geschrieben. Sollen alle Gruppen aktiviert werden, kann anstelle der Summe auch *reset* benutzt werden.

Gruppe	T1	T2	T3	T4	B1	B2	B3	B4	B5
Bit	0	1	2	3	4	5	6	7	8
Wert	1	2	4	8	16	32	64	128	256

Tab. 9.6.: Die Zuordnung der Masken-Bits und ihre Werte zu den einzelnen Szintillatorgruppen.

Verwendung des ILA

Um den ILA zu verwenden, wird der Befehl *trigger ila* benutzt. Das Steuerprogramm des ILA liegt auf der DAQ00. Wird ein Befehl auf der TRD01 eingegeben, wird mit dem Skript der entsprechende Befehl durch einen SSH-Tunnel auf der DAQ00 ausgeführt. Die Ausgabe erfolgt auf der TRD01. Die einzelnen Funktionen sind in Tabelle 9.7 gelistet. In Tabelle 9.8 sind die zur Zeit verfügbaren Trigger für den ILA aufgelistet. Mehrere Trigger können aktiviert werden, indem die Werte summiert werden.

Befehl	Beschreibung
read	Liest den Inhalt des RAM und gibt ihn auf der Konsole aus
ana	Liest den Inhalt des RAM und ruft das ROOT-Makro auf
arm	Aktiviert die ausgewählten Trigger
trg_mask	Auswahl der Trigger
info	Gibt verschiedene Informationen über den Status des ILA aus
mode	Wahl zwischen Pre- und Post-Modus
presamples	Die Anzahl der TTC-Takte vor dem Trigger, die nicht überschrieben werden sollen

Tab. 9.7.: Optionen des ILA

Bit	Wert	Bedingung
0	1	Manueller Trigger
1	2	Pretrigger wurde erzeugt
2	4	Verschiedene Trigger an GTU und Supermodul
3	8	Level 1 wurde generiert
4	16	Koinzidenz wurde gefunden
5	32	Level 1 Beitrag der GTU erhalten
6	64	Fallende Flanke vom GTU-busy
7	128	Level 0 wurde erzeugt
8	256	Bottom liefert ein Signal
9	512	Steigende Flanke GTU-busy
10	1024	Triggersequenz bei Noise gestoppt, weil Szintillator aktiv
11	2048	Fallende Flanke GTU-busy im Status <i>Idle</i>
12	4096	Level 1 überschrieben
13–15		Nicht vergeben

Tab. 9.8.: Die verfügbaren Trigger für den ILA. Um mehr als einen Trigger zu aktivieren, müssen die Werte addiert werden.

Einzelne Pretrigger Sequenzen

In Tabelle 9.9 sind die Optionen für den Trigger, mit dem einzelne Pretrigger erzeugt werden können, zusammengefasst. Die Konfiguration bleibt nach einer Sequenz erhalten.

Option	Beschreibung
amount	Anzahl der Pretrigger, die ausgegeben werden sollen
time	Zeit zwischen zwei Pretriggern. In TTC-Takten

Tab. 9.9.: Konfiguration für einzelne Pretrigger

Konfiguration der Triggereinheit ausgeben

Mit dem Befehl *trigger info* wird die komplette Konfiguration und der Zustand der Triggereinheit in eine Textdatei geschrieben. Dem Namen der Textdatei wird das aktuelle Datum und die Uhrzeit angehängen. Dadurch können die Triggereinstellungen beispielsweise für die Datennahme archiviert werden. Die Textdatei wird automatisch in das Verzeichnis kopiert, in dem der Befehl aufgerufen wird.

Zu den archivierten Werten gehören die Zähler für die Szintillatoren, Koinzidenz und ausgegebenen Triggern. Außerdem wird der aktuelle Trigger mit den Zusatzeinstellungen gespeichert. Zusätzlich werden noch die Einstellungen des Stresstests und die Zeitvorgaben der Zustandsmaschine gespeichert.

In diesem Kapitel wurde die Benutzerschnittstelle und ihre Optionen vorgestellt. Alle relevanten Überwachungswerte lassen sich über ein TCP/IP-Netzwerk mit einem DIM-Client beobachten. Zusätzlich werden einige ausgewählte Daten laufend im SMMon angezeigt und aktualisiert.

Um die Triggereinheit zu konfigurieren, wird der Command Channel des DIM-Clients benutzt. Dazu steht ein Skript auf dem Computer TRD01 zur Verfügung. Mit dem Skript kann außerdem der ILA kontrolliert werden.

Indem eine Benutzerschnittstelle entwickelt wurde, die sich komplett über ein TCP/IP-Netzwerk bedienen lässt, ist es möglich die Triggerauswahl und -konfiguration in Skripte für Tests einzubauen. Außerdem ist die Bedienung der elementaren Funktionen sehr einfach und erfordert keine oder nur eine kurze Einarbeitungszeit.

10. Zusammenfassung und Ausblick

ALICE ist eins von vier großen Experimenten am LHC (CERN). Ziel ist es, ein QGP zu untersuchen, das bei Pb–Pb–Kollisionen erzeugt wird. Ende 2010 fanden die ersten Pb–Pb–Kollisionen bei $\sqrt{s_{NN}} = 2,76$ TeV statt und erste Ergebnisse wurden bereits veröffentlicht. Aufgabe des TRD bei ALICE ist es, Elektronen mit hohem Transversalimpuls zu identifizieren und auf dieser Grundlage einen schnellen Triggerbeitrag zu liefern. Die einzelnen Supermodule des TRD werden im Institut für Kernphysik in Münster montiert. Um die Supermodule in Münster zu testen und zu kalibrieren, wurde im Rahmen dieser Arbeit ein Triggersystem entwickelt und getestet.

Für das System wurde ein vorhandener Aufbau aus [Bat07] übernommen und weiterentwickelt. Mit dem alten Aufbau war es zwar möglich die Spuren kosmischer Teilchen mit einem Supermodul aufzuzeichnen, allerdings brach die Stromversorgung in unregelmäßigen Abständen zusammen, da die Triggerrate kurzzeitig zu hoch war oder Triggersignale zu falschen Zeitpunkten an das Supermodul ausgegeben wurden.

Die neue Triggereinheit basiert auf einem Altera Cyclone FPGA, für den im Rahmen dieser Arbeit ein Design in VHDL entwickelt wurde. Die Triggereinheit verarbeitet die Signale von 90 Szintillatoren, die über und unter dem Supermodul platziert sind. Bei der Entwicklung der Triggereinheit hatte die Stabilität bei der Datennahme und den Tests Priorität. Deshalb werden die Trigger mit Hilfe einer Zustandsmaschine generiert. Zum ersten Mal ist es dabei möglich, vollständige Sequenzen bestehend aus einem Pretrigger, Level 0 und Level 1 an das Supermodul zu schicken. Dabei sind straffe Zeitvorgaben definiert, die garantieren, dass die Triggereinheit, das Supermodul und die Global Tracking Unit synchron laufen und die Triggerrate begrenzt wird. Seitdem die neue Triggereinheit eingesetzt wird, ist die Stromversorgung für die TRAPs nicht mehr zusammengebrochen.

Es stehen zwei verschiedene Trigger zur Verfügung, mit denen die Spuren von kosmischen Teilchen aufgezeichnet werden können. Beim Koinzidenztrigger wird das Supermodul ausgelesen, wenn ein kosmisches Teilchen die beiden Szintillatorebenen über und unter dem Supermodul durchquert hat. Beim GTU–Trigger nimmt die GTU eine Online–Rekonstruktion der Supermodul–Daten vor und gibt auf dieser Grundlage einen Level 1–Beitrag an die Triggereinheit zurück. Außerdem können die beiden Trigger kombiniert werden, so dass die Vorteile von beiden genutzt werden. Insgesamt wird dann eine Triggerrate von ca. 250 Signalen pro Sekunde für ein komplettes Supermodul erwartet. Die aufgenommenen Daten wurden teilweise bereits in [Alb10], [Gat10] und [Wal10] zur Kalibrierung benutzt.

Zusätzlich stehen weitere Trigger für diverse Tests zur Verfügung. Darunter fällt der Stresstest, bei dem zufällige Triggersequenzen mit beliebigen Raten erzeugt werden können. Dadurch kann die

Datennahme mit dem Supermodul mit hohen Raten simuliert werden. Der Noise-Trigger generiert Trigger, um das Rauschen der einzelnen Kanäle des Supermoduls aufzunehmen. Indem dabei die Szintillatorsignale als Veto verarbeitet werden, wird die Chance, dass zufällig kosmische Teilchen aufgezeichnet werden, erheblich minimiert.

Bei der Entwicklung der Triggereinheit wurde zusätzlich ein interner Logikanalysator entwickelt, mit dem sich bis zu 16 Signale in der Triggereinheit überwachen lassen. Die Daten können auf der Konsole oder mit einem ROOT-Skript ausgegeben werden.

Zur Kontrolle der Triggereinheit wurde eine Benutzerschnittstelle entwickelt, mit der es möglich ist, alle relevanten Parameter über ein TCP/IP-Netzwerk zu überwachen und zu konfigurieren. Damit wurde die Benutzerfreundlichkeit maßgeblich gesteigert und der Zugriff auf die Einheit durch Testprogramme oder -skripte ermöglicht. Dadurch können viele Prozesse automatisiert werden.

Die Entwicklung der Triggereinheit und der Aufbau des Triggersystems ist abgeschlossen und Erweiterungen sind zunächst nicht geplant. Durch die flexible Gestaltung der Triggereinheit ist es jedoch möglich, weitere Tests für das Supermodul zu entwickeln.

Ein sinnvoller Test für die Zukunft wäre z.B., die Daten des Supermoduls auf Bitfehler zu überprüfen. Dazu kann der Stresstest-Trigger benutzt werden und so konfiguriert werden, dass das GTU-busy berücksichtigt wird und die Daten tatsächlich ausgelesen werden. Indem die Daten mit hohen Raten auf Bitfehler geprüft werden, können Fehler bei der Datenübertragung entdeckt werden. Dazu wird am CERN der sogenannte *Online-Checker* bereits erfolgreich eingesetzt. Er muss lediglich für Münster angepasst werden.

Das Triggersystem wird bereits seit Supermodul IX erfolgreich eingesetzt. Bei Fertigstellung dieser Arbeit befindet sich gerade das Supermodul XII in der Montage in Münster. Insgesamt kann das neue Triggersystem somit für zehn von 18 Supermodulen verwendet werden.

Zusammenfassend ist zu vermerken, dass das neue Triggersystem sehr stabil funktioniert, einfach kontrolliert werden kann und diverse Trigger zur Aufzeichnung kosmischer Teilchen und für Tests des Supermoduls zur Verfügung stehen.

Anhang

A. Tabellen zum Triggersystem bei ALICE

	L0 Pb–Pb	L0 (pp)	L1 (Pb–Pb)	L2
1	V0 Minimum Bias	V0 Minimum Bias	TRD unlike e-pair high p_t	spare
2	V0 Semi-Central	V0 High Multiplicity	TRD like e-pair high p_t	spare
3	V0 Central	T0 Right	TRD jet low p_t	spare
4	T0 Multiplicity	T0 Left	TRD jet high p_t	spare
5	T0 Beam Gas	T0 Multiplicity	TRD electron	spare
6	PHOS MB	T0 Beam Gas	TRD hadron low p_t	spare
7	EMCAL MB	PHOS MB	TRD hadron high p_t	
8	Cosmic Telescope	EMCAL MB	ZDC 1	
9	DM like high p_t	Cosmic Telescope	ZDC 2	
10	DM unlike high p_t	DM like high p_t	ZDC 3	
11	DM like low p_t	DM unlike high p_t	ZDC special	
12	DM unlike low p_t	DM like low p_t	Topological 1	
13	DM single	DM unlike low p_t	Topological 2	
14	TRD pre-trigger	DM single	PHOS jet low p_t	
15	spare	TRD pre-trigger	PHOS jet high p_t	
16	spare	spare	EMCAL jet high p_t	
17	spare	spare	EMCAL jet med p_t	
18	spare	spare	EMCAL jet low p_t	
19	spare	spare	spare	
20	spare	spare	spare	
21	spare	spare		
22	spare	spare		
23	spare	spare		

Abb. A.1.: Die Triggereingänge der CTP. Für den Level 0 wird unterschieden, ob Proton–Proton– oder Blei–Blei–Kollisionen aufgenommen werden sollen. Die Level 2–Eingänge werden im Moment noch nicht benutzt. Bild aus [ALI03]

No.	Description	Condition
1.	MB	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [ZDC_1]_{L1}$
2.	SC	$[T0.V0_{SC} \cdot TRD_{pre}]_{L0} [ZDC_2]_{L1}$
3.	CE	$[T0.V0_{CE} \cdot TRD_{pre}]_{L0} [ZDC_3]_{L1}$
4.	$DM_{unlike,highp_t} \cdot TPC.MB$	$[T0.V0_{MB} \cdot DM_{unlike,highp_t} \cdot TRD_{pre}]_{L0} [ZDC_1]_{L1}$
5.	$DM_{unlike,highp_t} \cdot TPC.SC$	$[T0.V0_{SC} \cdot DM_{unlike,highp_t} \cdot TRD_{pre}]_{L0} [ZDC_2]_{L1}$
6.	$DM_{unlike,highp_t} \cdot No\ TPC.MB$	$[T0.V0_{MB} \cdot DM_{unlike,highp_t}]_{L0} [ZDC_1]_{L1}$
7.	$DM_{unlike,lowp_t} \cdot No\ TPC.MB$	$[T0.V0_{MB} \cdot DM_{unlike,lowp_t}]_{L0} [ZDC_1]_{L1}$
8.	$DM_{unlike,lowp_t} \cdot No\ TPC.SC$	$[T0.V0_{SC} \cdot DM_{unlike,lowp_t}]_{L0} [ZDC_2]_{L1}$
9.	$DM_{like,highp_t} \cdot TPC.MB$	$[T0.V0_{MB} \cdot DM_{like,highp_t} \cdot TRD_{pre}]_{L0} [ZDC_1]_{L1}$
10.	$DM_{like,highp_t} \cdot TPC.SC$	$[T0.V0_{SC} \cdot DM_{like,highp_t} \cdot TRD_{pre}]_{L0} [ZDC_2]_{L1}$
11.	$DM_{like,highp_t} \cdot No\ TPC.MB$	$[T0.V0_{MB} \cdot DM_{like,highp_t}]_{L0} [ZDC_1]_{L1}$
12.	$DM_{like,lowp_t} \cdot No\ TPC.MB$	$[T0.V0_{MB} \cdot DM_{like,lowp_t}]_{L0} [ZDC_1]_{L1}$
13.	$DM_{like,lowp_t} \cdot No\ TPC.SC$	$[T0.V0_{SC} \cdot DM_{like,lowp_t}]_{L0} [ZDC_2]_{L1}$
14.	$DM_{single} \cdot TRD_e.MB$	$[T0.V0_{MB} \cdot DM_{si} \cdot TRD_{pre}]_{L0} [TRD_e.ZDC_1]_{L1}$
15.	$DM_{single} \cdot TRD_e.SC$	$[T0.V0_{SC} \cdot DM_{si} \cdot TRD_{pre}]_{L0} [TRD_e.ZDC_2]_{L1}$
16.	$TRD_e.MB$	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [TRD_e.ZDC_1]_{L1}$
17.	$TRD_{lowp_t}.MB$	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [TRD_{lowp_t}.ZDC_1]_{L1}$
18.	$TRD_{highp_t}.MB$	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [TRD_{highp_t}.ZDC_1]_{L1}$
19.	$TRD_{unlike,highp_t}.MB$	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [TRD_{unlike,highp_t}.ZDC_1]_{L1}$
20.	$TRD_{unlike,highp_t}.SC$	$[T0.V0_{SC} \cdot TRD_{pre}]_{L0} [TRD_{unlike,highp_t}.ZDC_2]_{L1}$
21.	$TRD_{like,highp_t}.MB$	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [TRD_{like,highp_t}.ZDC_1]_{L1}$
22.	$TRD_{like,highp_t}.SC$	$[T0.V0_{SC} \cdot TRD_{pre}]_{L0} [TRD_{like,highp_t}.ZDC_2]_{L1}$
23.	$TRD_{jet,highp_t}.SC$	$[T0.V0_{SC} \cdot TRD_{pre}]_{L0} [TRD_{jet,highp_t}.ZDC_1]_{L1}$
24.	$TRD_{jet,lowp_t}.MB$	$[T0.V0_{MB} \cdot TRD_{pre}]_{L0} [TRD_{jet,lowp_t}.ZDC_1]_{L1}$
25.	$TRD_{jet,lowp_t}.SC$	$[T0.V0_{SC} \cdot TRD_{pre}]_{L0} [TRD_{jet,lowp_t}.ZDC_2]_{L1}$
26.	$PHOS_{highp_t}.MB$	$[T0.V0_{MB} \cdot PHOS_{highp_t} \cdot TRD_{pre}]_{L0} [ZDC_1]_{L1}$
27.	$PHOS_{lowp_t}.MB$	$[T0.V0_{MB} \cdot PHOS_{lowp_t} \cdot TRD_{pre}]_{L0} [ZDC_1]_{L1}$
28.	$PHOS_{lowp_t}.SC$	$[T0.V0_{SC} \cdot PHOS_{lowp_t} \cdot TRD_{pre}]_{L0} [ZDC_2]_{L1}$
29.	PHOS Stand-alone	$[T0.V0_{MB} \cdot PHOS_{MB}]_{L0} [ZDC_1]_{L1}$
30.	$EMCAL_{jet,highp_t}.MB$	$[T0.V0_{MB} \cdot EMCAL_{jet,highp_t}]_{L0} [ZDC_1]_{L1}$
31.	$EMCAL_{jet,medp_t}.MB$	$[T0.V0_{MB} \cdot EMCAL_{jet,medp_t}]_{L0} [ZDC_1]_{L1}$
32.	$EMCAL_{jet,lowp_t}.MB$	$[T0.V0_{MB} \cdot EMCAL_{jet,lowp_t}]_{L0} [ZDC_1]_{L1}$
33.	$EMCAL_{jet,lowp_t}.SC$	$[T0.V0_{SC} \cdot EMCAL_{jet,lowp_t}]_{L0} [ZDC_2]_{L1}$
34.	ZDC_{diss}	$[BX]_{L0} [ZDC_{spe}]_{L1}$
35.	cosmic	$[BX.cosmic_telescope]_{L0}$
36.	beam gas	$[T0_{beamgas}]_{L0}$

Abb. A.2.: Die Klassen, die in der CTP für Pb–Pb–Kollisionen definiert sind. In einer Klasse sind jeweils die Bedingungen um einen Level 0 bzw. Level 1 zu erzeugen zusammengefasst. Bild aus [ALI03]

B. Logikpegel

Für das Triggersystem in Münster werden verschiedene Logikpegel benutzt. Sie sind in Tab. B.1 zusammengefasst.

- NIM (*Nuclear Instrumentation Module*) ist ein Stromsignal, das mit $50\ \Omega$ terminiert wird.
- TTL (*Transistor-Transistor-Logik*) und CMOS 2,5V (*Complementary Metal Oxide Semiconductor*) sind Spannungssignale. Die Spannungsgrenzen sind für Ein- und Ausgänge separat definiert.
- LVDS (*Low Voltage Differential Signaling*) und ECL (*Emitter-Coupled Logic*) sind differentielle Signale. Das heißt, dass das Signal über zwei Leitungen übertragen wird. Der Logikpegel wird bestimmt, indem die Spannungsdifferenz zwischen den beiden Leitungen bestimmt wird. Bei einem Pegelwechsel werden die Leitungen umgepolt.

	Logisch 1	Logisch 0	Differentiell	V_{cc}/V_{ee}
NIM	-32 bis 12 mA -1,6 bis -0,6 V @ $50\ \Omega$	0 mA 0 V @ $50\ \Omega$	nein	5 V
TTL	EIN: $\geq 2,0\ \text{V}$ AUS: $\geq 2,4\ \text{V}$	EIN: $\leq 0,8\ \text{V}$ AUS: $\leq 0,4\ \text{V}$	nein	5 V
CMOS 2,5V	EIN: $\geq 1,7\ \text{V}$ AUS: $\geq 2,3\ \text{V}$	EIN: $\leq 0,7\ \text{V}$ AUS: $\leq 0,2\ \text{V}$	nein	2,5 V
LVDS	ΔU : +250 bis +400 mV	ΔU : -400 bis -250 mV	$V_{0L} = 1,0\ \text{V}$ $V_{0H} = 1,4\ \text{V}$	2,5 bis 3,5 V
ECL	ΔU : +800 mV	ΔU : -800 mV	$V_{0L} = -1,6\ \text{V}$ $V_{0H} = -0,8\ \text{V}$	-5,2 V

Tab. B.1.: Logikpegel verschiedener Standards. ΔU ist die Spannungsdifferenz je nach Polung der niedrigen Spannung V_{0L} und der hohen Spannung V_{0H} bei den differentiellen Standards. V_{cc} bzw. V_{ee} ist die benötigte Versorgungsspannung. Bei TTL und CMOS sind die Grenzen für EIN- bzw. AUS-gehende Signale verschieden.

C. Szintillatoren und Trefi-Module

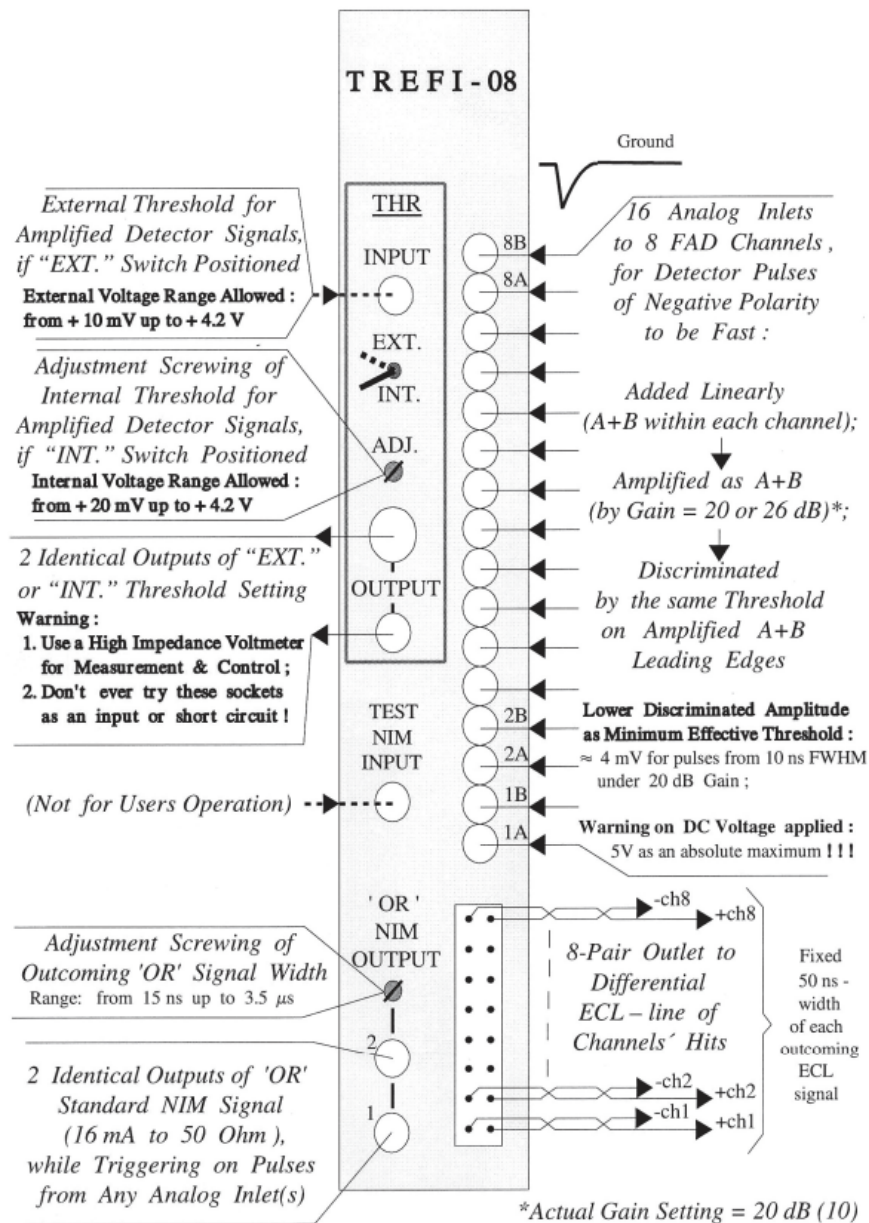


Abb. C.1.: Die Frontblende des TREFI-08-Moduls. Bild aus [Bat07]

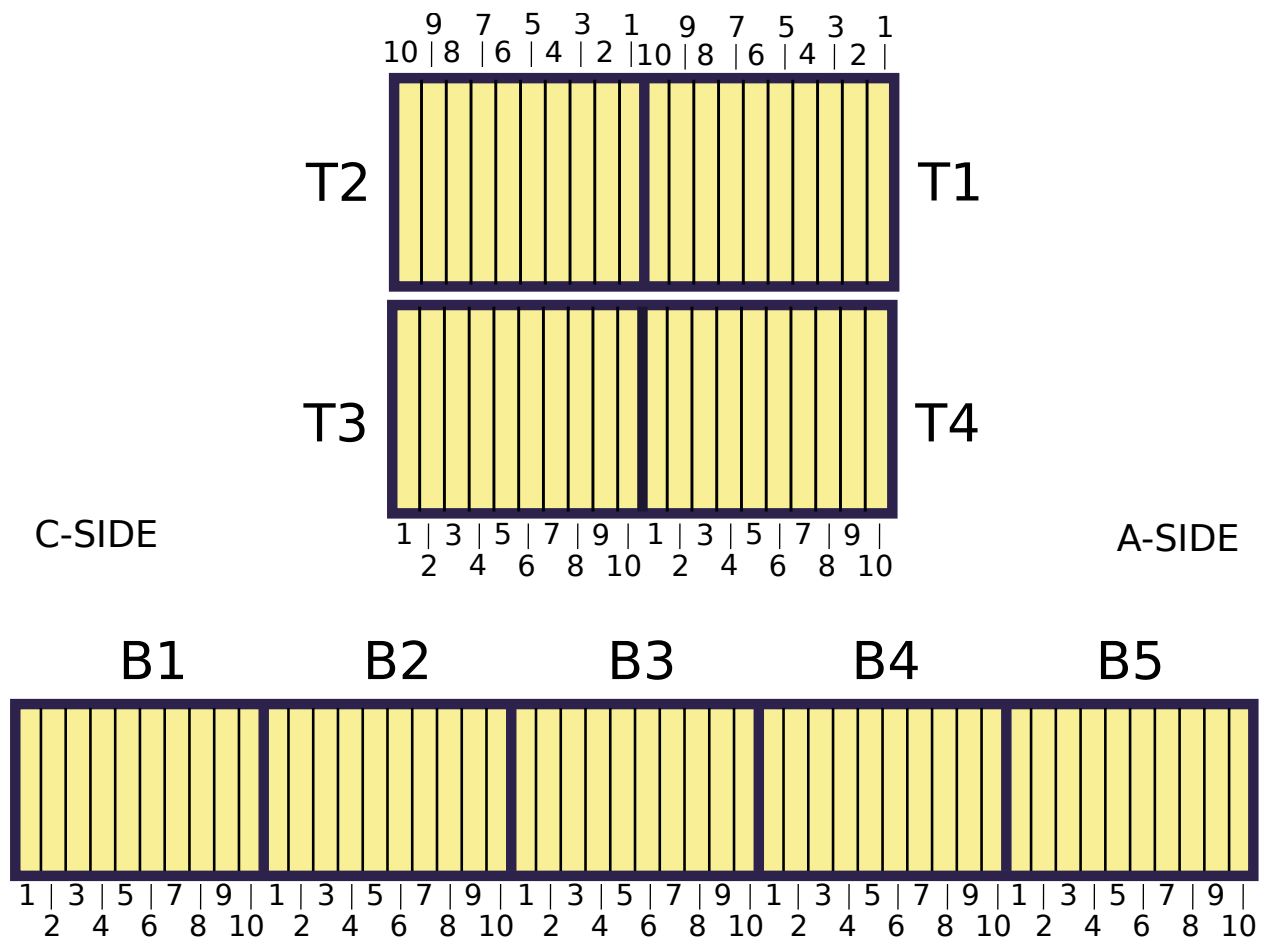
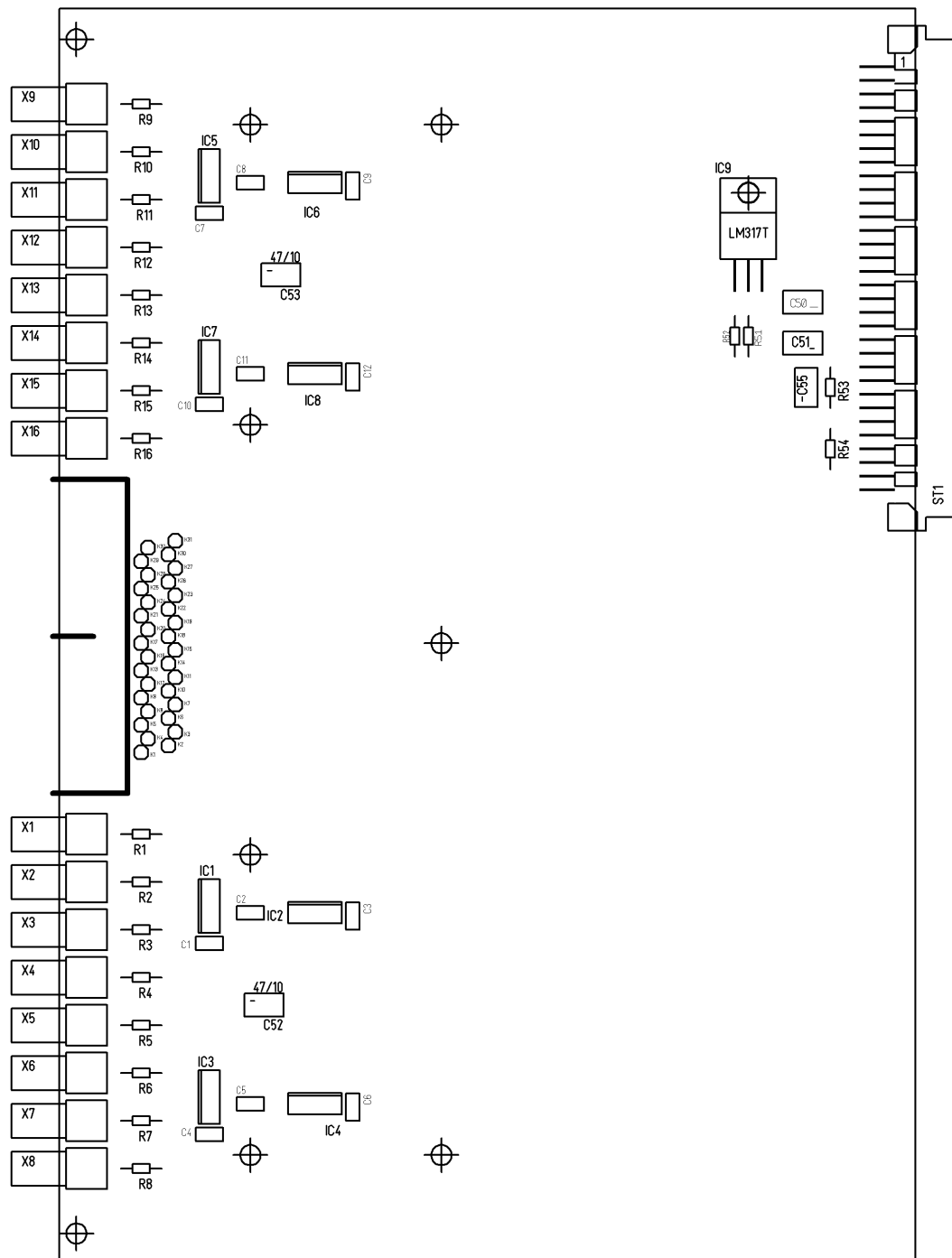


Abb. C.2.: Die Anordnung und Nummerierung der Szintillatoren in der oberen und unteren Ebene in der Draufsicht. In der Montagehalle ist die A-Side beim PC-Pool und die C-Side an der Außenwand der Halle.

D. Schaltpläne der NIM–LVDS– und LVDS–NIM–Wandler



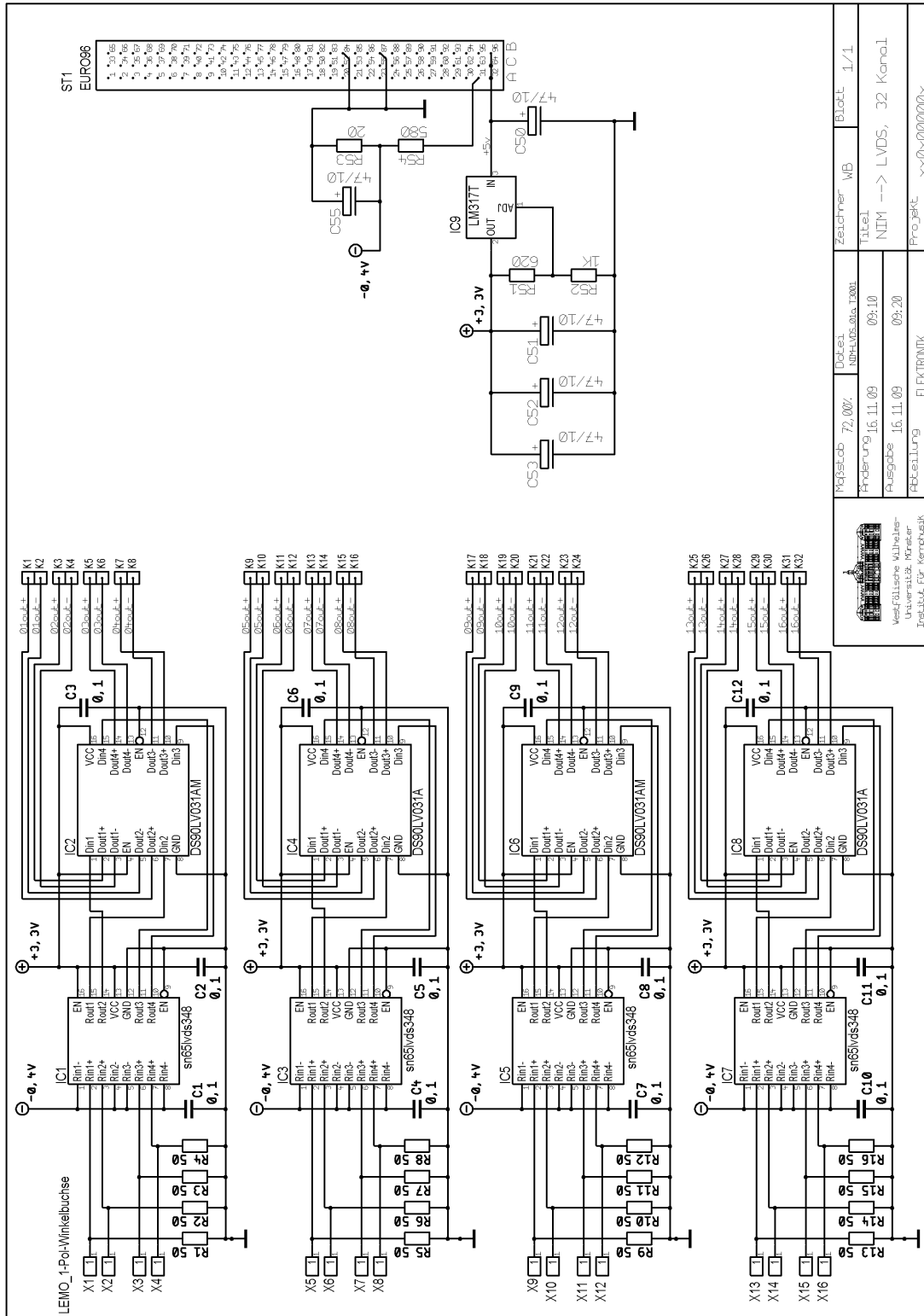
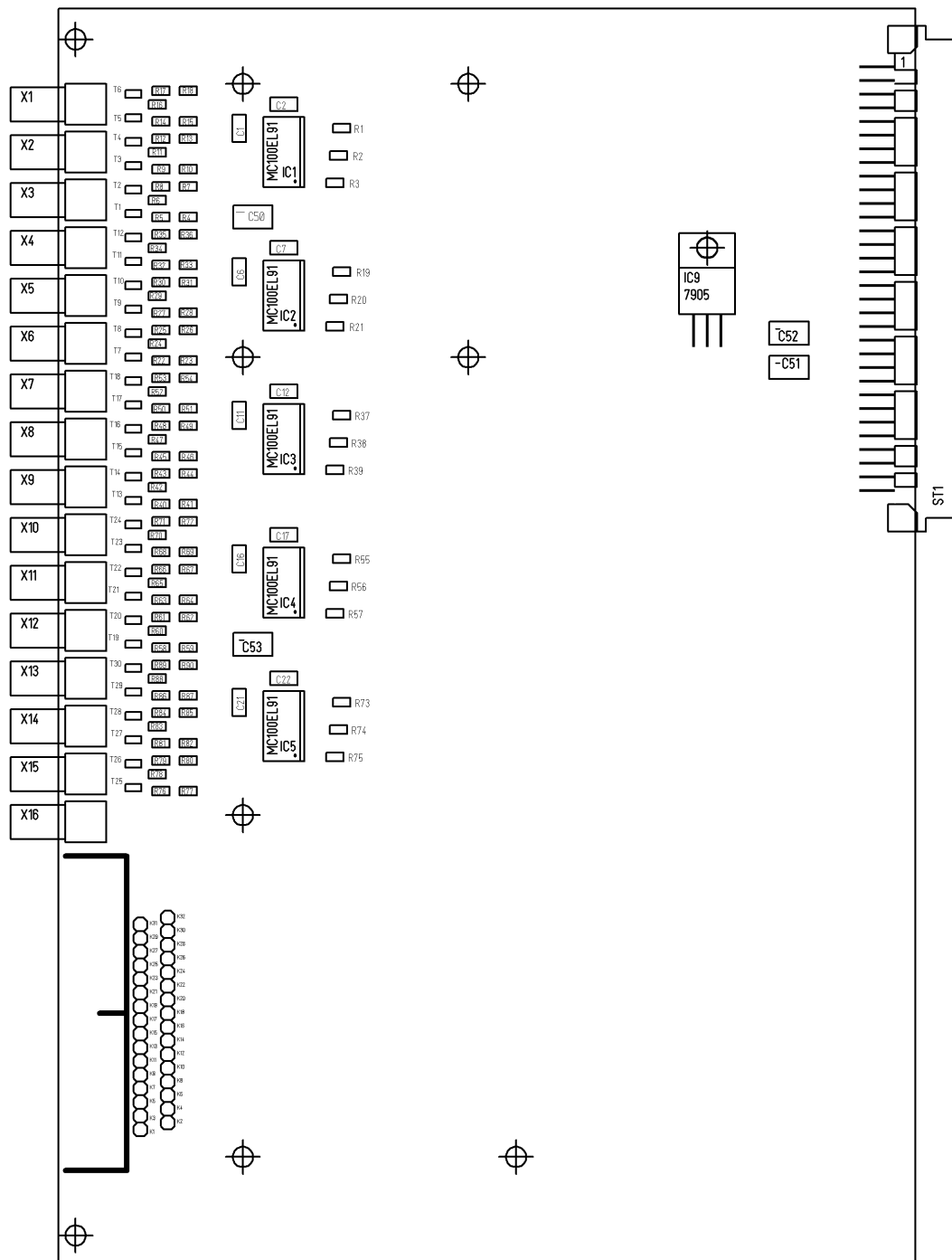


Abb. D.2.: Schaltplan NIM-LVDS-Wandler Schema



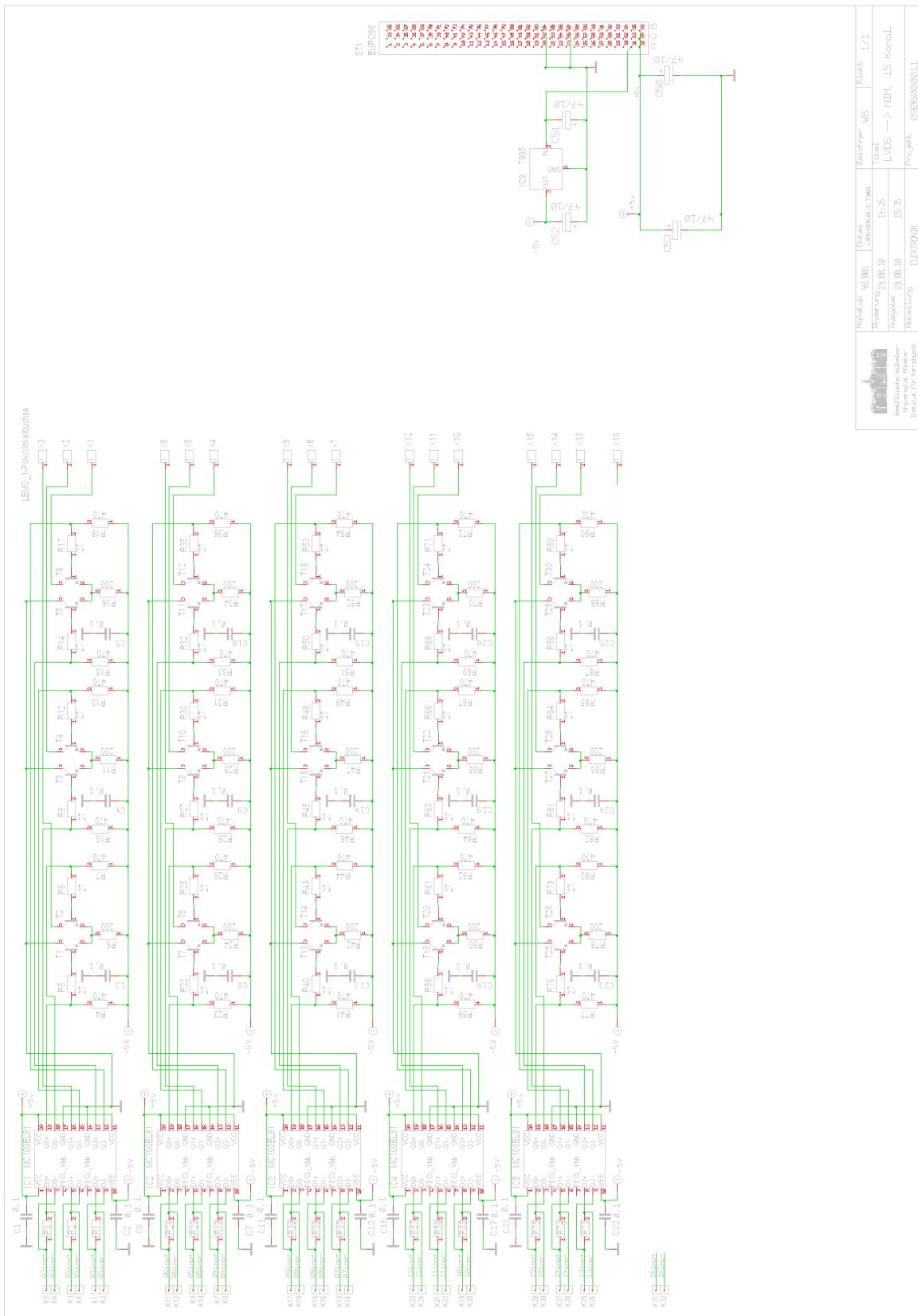
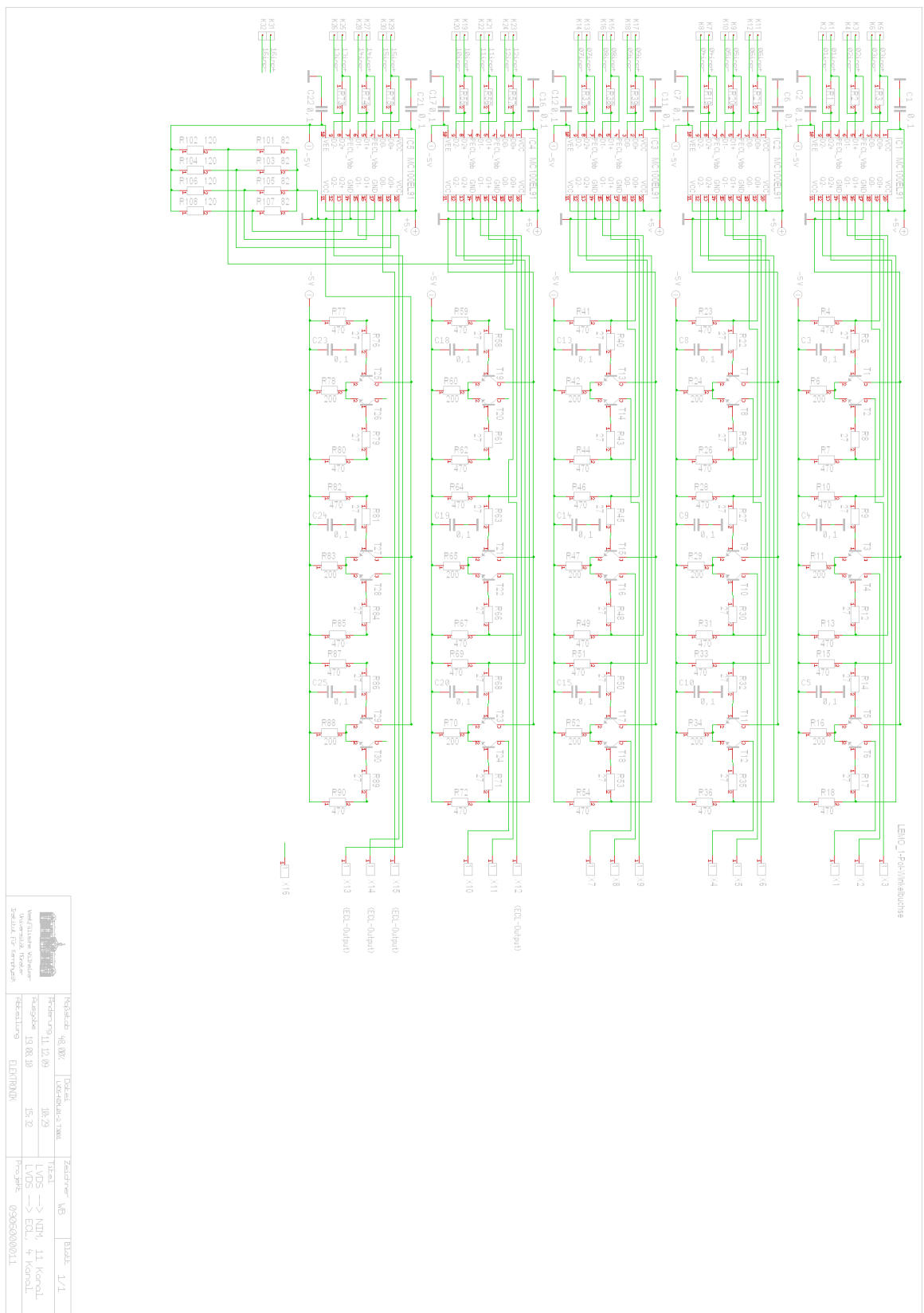


Abb. D.4.: Schaltplan LVDS–NIM–Wandler Schema 1

D. Schaltpläne der NIM–LVDS– und LVDS–NIM–Wandler



E. Register der Triggereinheit

Name	Adresse	Zugriff	Beschreibung
Szintillator-Zähler			
CNT_RST	0x0000	rw	Setzt Szintillator-Zähler zurück
CNT_ENA	0x0002	rw	Start/Stop Szintillator-Zähler
CNT_T1_H	0x0004	r	Zähler T1 – Bit 16–31
CNT_T1_L	0x0006	r	Zähler T1 – Bit 0–15
CNT_T2_H	0x0008	r	Zähler T2 – Bit 16–31
CNT_T2_L	0x000a	r	Zähler T2 – Bit 0–15
CNT_T3_H	0x000c	r	Zähler T3 – Bit 16–31
CNT_T3_L	0x000e	r	Zähler T3 – Bit 0–15
CNT_T4_H	0x0020	r	Zähler T4 – Bit 16–31
CNT_T4_L	0x0022	r	Zähler T4 – Bit 0–15
CNT_B1_H	0x0024	r	Zähler B1 – Bit 16–31
CNT_B1_L	0x0026	r	Zähler B1 – Bit 0–15
CNT_B2_H	0x0028	r	Zähler B2 – Bit 16–31
CNT_B2_L	0x002a	r	Zähler B2 – Bit 0–15
CNT_B3_H	0x002c	r	Zähler B3 – Bit 16–31
CNT_B3_L	0x002e	r	Zähler B3 – Bit 0–15
CNT_B4_H	0x0040	r	Zähler B4 – Bit 16–31
CNT_B4_L	0x0042	r	Zähler B4 – Bit 0–15
CNT_B5_H	0x0044	r	Zähler B5 – Bit 16–31
CNT_B5_L	0x0046	r	Zähler B5 – Bit 0–15
TOPR_H	0x0048	r	Zähler Top – Bit 16–31
TOPR_L	0x004a	r	Zähler Top – Bit 0–15
BTMR_H	0x004c	r	Zähler Bottom – Bit 16–31
BTMR_L	0x004e	r	Zähler Bottom – Bit 0–15
CNT_COINC_H	0x0060	r	Zähler Koinzidenz – Bit 16–31
CNT_COINC_L	0x0062	r	Zähler Koinzidenz – Bit 0–15
Diverse			
PDL_CTRL	0x0064	rw	Auswahl und Konfiguration der PDLs
Fortsetzung auf der nächsten Seite			

Tab. E.1 – Fortsetzung der letzten Seite

Name	Adresse	Zugriff	Beschreibung
PDL_DATA	0x0066	rw	Verzögerungswert für PDL
MODE	0x0068	rw	Triggerauswahl
MASK	0x006a	rw	Maske für Szintillatorgruppen
CNT_L0_LIMIT	0x200a	rw	Zeit zwischen Pretrigger und Level 0
CNT_L1_LIMIT	0x200c	rw	Zeit zwischen Pretrigger und Level 1
CNT_CLEAR_LIMIT	0x200e	rw	Timeout in Clear
CNT_READ_LIMIT	0x202e	rw	Zeit PT bis Ende Status Readout
Einzelne Triggersequenzen			
SNG_PT_REG	0x2004	w	Start-Bit für einzelne Sequenzen
SNG_PT_TIMES	0x2006	w	Anzahl der auszugebenden Pretrigger
SNG_PT_CNT_LIMIT	0x2008	w	Abstand zwischen den Pretriggern
Stresstest			
CNT_VF_LIMIT_H	0x2000	rw	Frequenz Pretrigger – Bit 16–31
CNT_VF_LIMIT_L	0x2002	rw	Frequenz Pretrigger – Bit 0–15
STRESS_L0_H	0x202a	rw	Wahrscheinlichkeit L0 – Bit 16–31
STRESS_L1_H	0x202c	rw	Wahrscheinlichkeit L1 – Bit 16–31
STRESS_L0_L	0x204a	rw	Wahrscheinlichkeit L0 – Bit 0–15
STRESS_L1_L	0x204c	rw	Wahrscheinlichkeit L1 – Bit 0–15
STRESS_PT_H	0x2040	rw	Wahrscheinlichkeit PT – Bit 16–31
STRESS_PT_L	0x2042	rw	Wahrscheinlichkeit PT – Bit 0–15
Trigger-Zähler			
CNT_PRE_H	0x2200	r	Pretrigger – Bit 16–31
CNT_PRE_L	0x2202	r	Pretrigger – Bit 0–15
CNT_L0_H	0x2204	r	Level 0 – Bit 16–31
CNT_L0_L	0x2206	r	Level 0 – Bit 0–15
CNT_L1_H	0x2208	r	Level 1 – Bit 16–31
CNT_L1_L	0x220a	r	Level 1 – Bit 0–15
CNT_L1_OVR_H	0x220e	r	Level 1 überschrieben – Bit 16–31
CNT_L1_OVR_L	0x2220	r	Level 1 überschrieben – Bit 0–15
CNT_L1_COINC_H	0x2222	r	Level 1 nur Koin. – Bit 16–31
CNT_L1_COINC_L	0x2224	r	Level 1 nur Koin. – Bit 0–15
CNT_TRG_CTRL	0x220c	w	Start/Stop/Reset Zähler
Interner Logikanalysator			
ILA_RDADDRESS_I	0x2260	rw	Leseadresse des RAM
ILA_WRADDRESS_I	0x2262	rw	Schreibadresse des RAM
Fortsetzung auf der nächsten Seite			

Tab. E.1 – Fortsetzung der letzten Seite

Name	Adresse	Zugriff	Beschreibung
ILA_Q_I	0x2264	rw	RAM – Ausgang
ILA_ARM_REG	0x2266	rw	Bit zur Triggeraktivierung
ILA_VME_DONE_REG	0x2268	r	Nach VME–Auslese high
ILA_MODE	0x226a	rw	Pre oder Post Mode
ILA_CNT_LIMIT	0x226c	rw	Presamples einstellen
ILA_COMPLETE	0x226e	rw	Nach erfolgreichem Trigger high
ILA_TRG_MASK	0x2280	rw	Triggermaske
ILA_TRG_REG	0x2282	rw	Manueller Trigger
Zähler für Trefi–Eingänge			
CNT_TREFI0	0x2404	r	Trefi 0–Eingang
CNT_TREFI1	0x2406	r	Trefi 1–Eingang
CNT_TREFI2	0x2408	r	Trefi 2–Eingang
CNT_TREFI3	0x240a	r	Trefi 3–Eingang
CNT_TREFI4	0x240c	r	Trefi 4–Eingang
CNT_TREFI5	0x240e	r	Trefi 5–Eingang
CNT_TREFI6	0x2420	r	Trefi 6–Eingang
CNT_TREFI7	0x2422	r	Trefi 7–Eingang
CNT_TREFI8	0x2424	r	Trefi 8–Eingang
CNT_TREFI9	0x2426	r	Trefi 9–Eingang
CNT_TREFI_ALL	0x242c	r	Summe aller Eingänge
TREFI_BTMTOP_SELECT	0x2428	w	Top oder Bottom mit PDL
CNT_TREFI_CTRL	0x242a	w	Start/Stop/Reset Zähler

Tab. E.1.: Verfügbare Register in der Triggereinheit. Zusätzlich sind die Adressen und die erlaubten Zugriffe über den VME–Bus angegeben.

Abbildungsverzeichnis

2.1. a) Kopplungskonstante der starken Wechselwirkung b) Quark–Antiquark–Potential in Abhängigkeit des Abstandes	6
2.2. Phasendiagramm Quark–Gluon Plasma	7
2.3. Simulation einer Pb–Pb–Simulation bei $\sqrt{s_{NN}} = 5,5$ TeV	9
2.4. Energieverlust eines Myons in Kupfer pro Längeneinheit	12
2.5. Das Funktionsprinzip einer Vieldrahtkammer	13
2.6. Das Funktionsprinzip eines Photovervielfachers	15
3.1. Schematischer Aufbau des LHC	18
3.2. Vorbeschleuniger des LHC	20
3.3. Der ALICE–Detektor mit den Subdetektoren	21
3.4. Rekonstruktion einer der ersten Pb–Pb–Kollisionen bei ALICE	23
3.5. Ablauf einer ALICE–Triggersequenz	24
3.6. Das Triggersystem von ALICE schematisch	27
3.7. Das Datenformat des TTC B–Kanal	28
4.1. Aufbau des TRD – Segmentierung in Module	33
4.2. Querschnitt durch ein TRD–Modul	34
4.3. Funktionsweise des TRD schematisch	36
4.4. Signalformen von Elektronen und Pionen im TRD	38
4.5. Triggersystem des TRD schematisch	39
4.6. Zustandsmaschine in Kontrollbox B	41
4.7. Signalverarbeitung der Detektorelektronik des TRD schematisch	42
4.8. Zeitlicher Verlauf einer Triggersequenz	42
4.9. Aufbau der GTU schematisch	44
5.1. Vertikaler Fluss kosmischer Teilchen gegen die Höhe	49
6.1. Schematischer Aufbau zur Detektion kosmischer Teilchen in Münster	53
6.2. Signalverarbeitung der Szintillatoren	54
6.3. Schaltbild der Photovervielfacher	55
6.4. Blockdiagramm des CAEN V1495	57
6.5. Aufbau um Koinzidenzkurve aufzunehmen	59

6.6.	Koinzidenzkurven vor und nach Anpassen der Kabellängen	61
6.7.	Aufbau zur Kalibrierung der Trefi-Diskriminatorschwellen	62
6.8.	Koinzidenzkurven zur Kalibrierung der Diskriminatorschwellen	63
6.9.	Koinzidenzen zwischen Szintillator B2-1 und Top	65
6.10.	Ergebnisse für Diskriminatorschwellen für T1-T4	66
6.11.	Ergebnisse für Diskriminatorschwellen für B1-B4	67
6.12.	Ergebnisse für Diskriminatorschwelle für B5	68
6.13.	Zählraten der einzelnen Szintillatorgruppen	69
6.14.	Zählraten der unteren und oberen Szintillatorebene	69
7.1.	Schematischer Aufbau des Triggersystems in Münster	74
7.2.	Foto des elektrischen Aufbaus in Münster	75
7.3.	Synchronisation der Triggersignale am TTCex-Eingang	75
7.4.	Die Zustandsmaschine der TRAPs	78
7.5.	Die Zustandsmaschine der Triggereinheit	79
7.6.	Schematische Darstellung des RAM im FPGA	82
7.7.	Ausschnitt der Ausgabe des ILA auf der Konsole	84
7.8.	Bild des ROOT-Makros das die ILA-Daten auswertet	85
8.1.	Logik um einen Pretrigger mit der Koinzidenzbedingung zu erzeugen	89
8.2.	Koinzidenz-Triggersequenz mit ILA aufgenommen	90
8.3.	Vergößerung einer Sequenz um Pretrigger bei Koinzidenz	91
8.4.	Koinzidenzrate in Abhängigkeit von der Top- und Bottom-Pulslänge	92
8.5.	GTU-Trigger-Sequenz mit dem ILA aufgenommen	93
8.6.	Rauschen der einzelnen Kanäle von Lage vier des Supermoduls XI	101
8.7.	Erzeugung von Zufallszahlen mit einem LFSR	102
8.8.	Verhalten der Zustandsmaschine beim Stresstest	103
8.9.	Aufnahme des ILAs vom Stresstest	104
9.1.	Zugriff auf die Triggereinheit schematisch	110
9.2.	Funktionsprinzip des Distributed Information Management System	111
9.3.	DIM-Client, der auf den Server der Triggereinheit zugreift	112
9.4.	SMMon-Modul, das die Informationen der Triggereinheit darstellt	113
A.1.	Die Triggereingänge der CTP	123
A.2.	Die Klassen in der CTP für Pb-Pb-Kollisionen	124
C.1.	Die Frontblende des TREFI-08-Moduls	127
C.2.	Anordnung und Nummerierung der Szintillatoren in Münster	128
D.1.	Schaltplan NIM-LVDS-Wandler Platine	130

D.2. Schaltplan NIM–LVDS–Wandler Schema	131
D.3. Schaltplan LVDS–NIM–Wandler Platine	132
D.4. Schaltplan LVDS–NIM–Wandler Schema 1	133
D.5. Schaltplan LVDS–NIM–Wandler Schema 2	134

Tabellenverzeichnis

2.1. Quarks und Leptonen im Standardmodell	4
2.2. Wechselwirkungen im Standardmodell und die Gravitation	6
6.1. Diskriminatorschwellen der einzelnen Szintillatorgruppen	65
6.2. Zählraten der einzelnen Szintillatorgruppen	68
6.3. Zählraten der Szintillatorebenen	70
7.1. Signallaufzeiten des Pretriggers	77
7.2. Die Zeitvorgaben der Zustandsmaschinen der TRAPs und der Triggereinheit	80
7.3. Beschreibung der einzelnen Signale, die vom ILA überwacht werden	84
8.1. Generierte Level 1 im kombinierten Betrieb Koinzidenz- und GTU-Trigger	97
8.2. Die einzelnen Raten beim kombinierten Trigger	98
8.3. Zählraten der einzelnen Trigger beim Noise-Trigger	100
8.4. Bedingungen unter denen Trigger beim Stresstest erzeugt werden	103
8.5. Verschiedene Einstellungen und Triggerraten beim Stresstest	105
9.1. Die Optionen, die beim Skript <i>trigger</i> zur Verfügung stehen	113
9.2. Verfügbare Trigger des Triggersystems	114
9.3. Zusätzliche Einstellungen für Trigger	114
9.4. Optionen, um die Zeitvorgaben der Zustandsmaschine anzupassen	115
9.5. Optionen, um den Stresstest-Trigger zu konfigurieren	115
9.6. Zuordnung der Masken-Bits zu den einzelnen Szintillatorgruppen	115
9.7. Optionen des ILA	116
9.8. Verfügbare Trigger für den ILA	116
9.9. Konfiguration für einzelne Pretrigger	117
B.1. Logikpegel verschiedener Standards	125
E.1. Verfügbare Register in der Triggereinheit	137

Literaturverzeichnis

- [Ada03] J. Adams et al: *Evidence from $d + Au$ Measurements for Final-State Suppression of High- p_t Hadrons in $Au + Au$ Collisions at RHIC*
Physical Review Letters, Vol. 91, 7, 72304 (2003)
- [Alb10] Diplomarbeit B. Albrecht: *Gain Calibration of ALICE TRD Modules Using Cosmic Ray Data*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2010
- [ALI01] ALICE Collaboration: *Transistion Radiation Detector Technical Design Report*
CERN/LHCC 2001-021
- [ALI03] ALICE Collaboration: *ALICE Technical Design Report of the Trigger, Data Acquisition, High-Level Trigger and Control System*
CERN-LHCC-2003-022
- [ALI04] ALICE Collaboration: *ALICE: Physics Performance Report, Volume 1*
J. Phys. G: Nucl. Part. Phys. 30 1517-1764 (2004)
- [ALI08] ALICE Collaboration: *Electromagnetic Calorimeter Technical Design Report*
CERN/LHCC 2008-014
- [ALI08a] ALICE Collaboration: *The ALICE experiment at the CERN LHC*
: J. Instrum. 3 (2008) S08002
- [ALI10] Homepage von ALICE: <http://aliceinfo.cern.ch/>, aufgerufen: 8.11.2010
- [ALI10a] ALICE Collaboration: *Charged-particle multiplicity density at mid-rapidity in central Pb-Pb collisions at $\sqrt{s_{NN}} = 2.76$ TeV*
arXiv:1011.3916
- [Alt03] Altera: *Cyclone FPGA Family Data Sheet*
San Jose, USA, 2003
- [Alt08] Altera: *Cyclone Device Handbook, Volume 1*
San Jose, USA, 2008

- [Ams07] C. Amsler: *Kern- und Teilchenphysik*
vdf Hochschulverlag AG an der ETH Zürich, Zürich, 2007
- [Ams08] C. Amsler et al.: *Cosmic Rays*
Physics Letters B667, 1 2008
- [Ang05] V. Angelov et al.: *ALICE TRAP User Manual*
Ruprecht–Karls–Universität Heidelberg (2005)
- [Ang08] V. Angelov: *Fast Data Processing in ALICE TRD FEE & GTU*
Vortrag ALICE Workshop, Sibiu, 2008
- [Ang10] V. Angelov: Persönliche Kommunikation, 06.03.2010
- [Ang11] V. Angelov, Physikalisches Institut, Ruprecht–Karls–Universität Heidelberg
- [Bat07] Diplomarbeit B. Bathen: *Aufbau eines Triggers für Tests der ALICE–TRD–Supermodule mit kosmischer Strahlung*
Institut für Kernphysik, Westfälische Wilhelms–Universität Münster, 2007
- [Ber02] C. Berger: *Elementarteilchenphysik*
Springer, Berlin Heidelberg New York, 2002
- [Brü04] O. S. Brüning et al.: *LHC Design Report Vol. 1*
CERN-2004-003-V-1
- [BMS07] P. Braun-Munzinger, J. Stachel: *The quest for the quark-gluon plasma*
Nature 448: 302-309 (2007)
- [CAE10] C.A.E.N.: *Technical Information Manual V1495 General Purpose Board*
Viareggio Solingen Staten Island, 2010
- [CER08] Communication Group: *CERN-Brochure-2008-001-Eng*, CERN 2008
- [CER10] LHC Machine Outreach: <http://lhc-machine-outreach.web.cern.ch/>, aufgerufen: 05.11.2010
- [Chr04] J. Christiansen et al.: *TTCrx Reference Manual*
CERN, Genf, 2004
- [Cuv03] Diplomarbeit J. de Cuveland: *Entwicklung der Spurrekonstruktionseinheit für den ALICE–Übergangsstrahlungsdetektor am LHC (CERN)*
Ruprecht–Karls–Universität Heidelberg, 2003
- [Dem10] W. Demtröder: *Experimentalphysik 4*
Springer, Berlin Heidelberg, 2010

- [Die10] T. Dietel et al.: *TRD Documentation Project*
<http://www.physi.uni-heidelberg.de/~alice/tdp/latest>, 2010
- [Die11] T. Dietel, Institut für Kernphysik, Westfälische Wilhelms-Universität Münster
- [DIM11] Homepage des DIM-Projekts: <http://dim.web.cern.ch/>, aufgerufen Januar 2011
- [Ems09] Dissertation D. Emschermann: *Construction and Performance of the ALICE Transition Radiation Detector*
Ruprecht-Karls-Universität Heidelberg, 2009
- [Eva04] D. Evans et al.: *ALICE Trigger System*
10th Workshop on Electronics for LHC and Future Experiments, Boston, MA, USA, 13 - 17 Sep 2004, pp.277-280
- [Gat10] Diplomarbeit H. Gatz: *Calibration and Alignment of ALICE TRD Super Modules Using Cosmic Ray Data*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2010
- [Gru00] C. Grupen: *Astroteilchenphysik: Das Universum im Licht der kosmischen Strahlung*
Vieweg, Braunschweig Wiesbaden, 2000
- [Jov04] P. Jovanovic: *Central Trigger Processor – Preliminary Design Review*,
<http://cdsweb.cern.ch/record/478892/files/p318.pdf>, 2004
- [Kal11] M. Kalisky, Institut für Kernphysik, Westfälische Wilhelms-Universität Münster
- [Kir07] Diplomarbeit S. Kirsch: *Development of the Supermodule Unit for the ALICE Transition Radiation Detector at the LHC (CERN)*
Ruprecht-Karls-Universität Heidelberg, 2007
- [Kle08] J. Klein: *Commissioning of and Preparations for Physics with the Transition Radiation Detector in A Large Ion Collider Experiment at CERN*
Ruprecht-Karls-Universität Heidelberg, 2008
- [Kwe09] MinJung Kweon for the ALICE Collaboration: *The Transition Radiation Detector for ALICE at LHC*
Nucl. Phys. A 830 (2009) 535c–538c
- [LAN10] Homepage Los Alamos National Laboratory: <http://www.lanl.gov/milagro/cosmicrays.shtml>,
aufgerufen 29.11.2010
- [NA00] N50 Collaboration: *Evidence of deconfinement of quarks and gluons from J/ψ suppression pattern measured in Pb-Pb collisions at CERN-SPS*
CERN-EP – 2000-013

- [PDG10] K Nakamura et al. (Particle Data Group) 2010 J. Phys. G: Nucl. Part. Phys. 37 075021
- [Pov06] B. Povh, Rith, Scholz, Zetsche: *Teilchen und Kerne*
Springer, Berlin Heidelberg New York, 2006
- [Pri10] Homepage des Departement of Computer Science, Princeton University:
<http://www.cs.princeton.edu>, aufgerufen: 13.01.2011
- [PTS10] Homepage des TRD Pretrigger System:
<http://alice.physi.uni-heidelberg.de/oyama/PreTrigger/pkww/>, aufgerufen: 11.11.2010
- [Sic09] Diplomarbeit E. Sicking: *Alignment of ALICE TRD Modules Using Cosmic Ray Data*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2009
- [Tay] B.G. Taylor: *TTC laser transmitter (TTCex, TTXtx, TTCmx) User Manual*
ttc.web.cern.ch/TTC/TTCtxManual.pdf
- [TM04] P. A. Tipler, G. Mosca: *Physik Für Wissenschaftler und Ingenieure*
Elsevier, München, 2006
- [Wal10] Diplomarbeit M. Walter: *Performance of Online Tracklet Reconstruction with the ALICE TRD*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2010
- [Web10] H. Weber: *UrQMD Animations*
<http://th.physik.uni-frankfurt.de/~weber/CERNmovies/index.html>, aufgerufen: 12.10.2010
- [WW06] C. Weinheimer, J.P. Wessels: *Skript zur Vorlesung Kern- und Teilchenphysik*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2006
- [Wil04] Diplomarbeit A. Wilk: *Elektronen-Pionen-Seperation im ALICE TRD*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2004
- [Wil10] Diplomarbeit A. Wilk: *Particle Identification Using Artificial Neural Networks with the ALICE Transition Radiation Detector*
Institut für Kernphysik, Westfälische Wilhelms-Universität Münster, 2004
- [YHM05] K. Yagi, T. Hatsuda, Y. Miake: *Quark-Gluon Plasma: From Big Bang to Little Bang*
Cambridge University Press, Cambridge New York Melbourne, 2005
- [Xil96] Xilinx, Peter Alfke: *Application Note: Efficient Shift Registers, LFSR Counters, and Long Pseudo-Random Sequence Generators*, San Jose, 1996

Danksagung

An dieser Stelle möchte ich allen Leuten danken, die maßgeblich zum guten Gelingen dieser Arbeit beigetragen haben. Zunächst ist hier Herr Prof. Wessels zu nennen, der mir das Anfertigen der Arbeit mit vielen Freiheiten und einen spannenden Aufenthalt am CERN ermöglicht hat.

Besonderer Dank gilt Dr. Thomas Dietel, der mich bei Fragen und Problemen stets unterstützt hat, oft hilfreiche Tipps geben konnte und viele Impulse zur Weiterentwicklung der Triggereinheit gegeben hat. Des Weiteren bedanke ich mich für die Korrektur der Arbeit.

Für die netten Hilfestellungen und die Entwicklung der Pegelwandler bedanke ich mich bei der elektronischen Werkstatt um Herrn Berendes. Besonders zu nennen ist dabei Herr Buglak.

Des Weiteren bedanke ich mich bei der ALICE-Kollaboration für die gute Zusammenarbeit. Besonders hervorzuheben sind hier Stefan Kirsch und Felix Rettig, die meine Fragen zur GTU stets geduldig beantwortet haben. Außerdem bedanke ich mich für die Hilfestellungen bei Problemen mit VHDL.

Bei der gesamten Arbeitsgruppe und insbesondere bei meinen Bürokollegen Andrea Nustede, Michael Kowalik, Sebastian Klamor und Stefan Korsten bedanke ich mich für die gute Atmosphäre während der Diplomarbeit. Sie hatten immer ein offenes Ohr für Schwierigkeiten bei ROOT und Latex.

Für die Korrektur der Arbeit bedanke ich mich bei Bastian Drees und Daniel Uhlmann. Bastian danke ich außerdem für die hilfreichen Gespräche während des Studiums. Außerdem bedanke ich mich bei Gabriel Pelegrina Bonilla, mit dem ich viele Hürden des Studiums gemeinsam bewältigt habe.

Für die großartige Zeit in Münster und die gelungenen Ablenkungen bedanke ich mich bei all meinen (ehemaligen) Mitbewohnern und Freunden.

Zum Schluss danke ich meiner Familie: Mechthild und Peter, die mir das Studium erst ermöglicht und mich stets in allen Lebenslagen unterstützt haben und Denis, der immer verfügbar war, um mir zu helfen und diese Arbeit Korrektur gelesen hat.

Ich versichere, dass ich die Arbeit selbstständig verfasst und keine anderen als die angegebenen Quellen und Hilfsmittel benutzt, sowie Zitate kenntlich gemacht habe.

Münster, Februar 2011

Jonas Anielski

Referent: Prof. Dr. J. P. Wessels

Korreferent: Prof. Dr. C. Weinheimer

