

Melanie Hoppe

Aufbau und Inbetriebnahme einer Funkenkammer

— 2004 —

Experimentelle Physik

Aufbau und Inbetriebnahme einer Funkenkammer

Diplomarbeit

von

Melanie Hoppe

Westfälische Wilhelms-Universität Münster

Institut für Kernphysik

— 2004 —

“...you know the sparks begin to fly”

Bob Dylan

für Christian

Inhaltsverzeichnis

Einleitung	3
1 Theoretische Grundlagen der Teilchenphysik	7
1.1 Der Aufbau der Materie	8
1.1.1 Quarks und Leptonen	8
1.1.2 Wechselwirkungen	9
1.1.3 Symmetrien und Erhaltungssätze	10
1.2 Wechselwirkung von Strahlung mit Materie	11
2 Kosmische Strahlung	15
2.1 Die primäre kosmische Strahlung	15
2.2 Wechselwirkung mit der Erdatmosphäre	18
2.3 Das Myon	20
3 Das Grundprinzip der Funkenkammer	23
3.1 Ionisierung eines Gases	24
3.2 Diffusion in feldfreiem Gas	28
3.3 Der Gasaustausch mit der Umgebung	30
3.4 Die Gasversorgung	32
3.5 Messung der Leckrate	33
4 Der Trigger	35
4.1 Der Szintillator und Lichtleiter	35
4.1.1 Die Funktionsweise von Szintillatoren	35
4.1.2 Der Aufbau der Paddles	36
4.2 Der Photomultiplier	38
4.2.1 Der Sekundärelektronenvervielfacher (SEV)	40

4.2.2	Beschaltung der Base	41
4.2.3	Pulsformmessung der Photomultiplier	42
4.3	Die Triggerelektronik	45
4.4	Die geometrische Akzeptanz und Effizienz des Triggers	48
4.4.1	Die geometrische Akzeptanz der Triggerdetektoren	49
4.4.2	Die Effizienz der Triggerdetektoren	49
4.4.3	Die Koinzidenzrate	51
5	Das Zündsystem	57
5.1	Die Zündschaltung	57
5.1.1	Die Zündkerze	58
5.2	Die Zündansteuerung	59
5.2.1	Der Gleichrichter	60
5.2.2	Der Elektrolytkondensator	61
5.2.3	Der Triac	61
5.2.4	Abschlusswiderstand	62
5.3	Aufbau der Zündschaltung und -ansteuerung	62
5.4	Optimierung der Funkenkammer für den Dauerbetrieb	65
A	Die technischen Daten der Funkenkammer	69
B	Die Konstruktionspläne	73
	Literaturverzeichnis	87
	Danksagung	89

Einleitung

Eine Funkenkammer ist ein optischer Detektor für ionisierende Strahlung. Mit einer auf dem Erdboden aufgestellten Funkenkammer werden beispielsweise Spuren kosmischer Teilchen sichtbar. Die kosmische Strahlung und deren in der Erdatmosphäre entstehende sekundäre Komponente ist mit einem Anteil von zwölf Prozent eine der Quellen der natürlichen Radioaktivität. Die für den Menschen nicht wahrnehmbaren Teilchen hinterlassen in der Kammer Spuren, die als Abfolge von Funken beobachtet werden können. Das Thema dieser Arbeit ist der Aufbau und die Inbetriebnahme einer solchen Funkenkammer zu Demonstrationszwecken.

Funkenkammern stehen in der Tradition der optischen Spurendetektoren, die seit mehr als 100 Jahren unersetzbare Dienste bei der Entdeckung neuer Teilchen leisten. Die Messung erfolgt über die Wechselwirkung des Teilchens mit der Materie des Detektors, wobei das Detektormaterial selbst als Target dienen kann. Als erste *Teilchendetektoren* sind *Emulsionen* zu nennen. So entdeckte A. H. Becquerel Ende des 19. Jahrhunderts die Radioaktivität durch das Schwärzen von Fotoplaten.

Die *Nebelkammer* wurde 1912 von C. T. R. Wilson erfunden und hat zahllose wichtige Erkenntnisse über das Verhalten von Elementarteilchen erbracht. So erhielt C. D. Anderson 1936 für die Entdeckung der Antimaterie in Form des Positrons mit der Nebelkammer den Nobelpreis [And36]. Die Nebelkammer besteht aus einem Gefäß, das mit einem wasserdampfgesättigten Gas gefüllt ist und dessen Volumen durch einen verschiebbaren Kolben expandiert werden kann. Die adiabatische Ausdehnung kühlt das Gas ab und es ist nun mit Wasserdampf übersättigt. Dieser scheidet sich aber erst dann in Form kleiner Wassertropfen (Nebel) ab, wenn Kondensationskeime, zum Beispiel kleine Staubeilchen oder auch geladene Gasmoleküle, vorhanden sind. Wenn sich unmittelbar nach der Expansion geladene Teilchen durch das Gas bewegen, ionisieren sie auf ihrem Weg Atome und Moleküle. An den entstandenen Ionen kondensiert der Wasserdampf zu Tröpfchen, ehe sich die Ionen durch Diffusion merklich verschieben können. Wird gleichzeitig von der Seite beleuchtet, so wird das Licht an diesen Tröpfchen gestreut und es ist möglich, die Spur des Teilchens als Nebelstreifen auf dem dunklen Untergrund des geschwärzten Kolbens zu beobachten oder zu fotografieren.

Die Nebelkammer hat den Nachteil, dass schwach ionisierende Teilchen im Füllgas wegen dessen geringer Dichte zu wenig Ionen erzeugen. Die Entwicklung nachfolgender Detektoren mit höherer Targetdichte versprach eine Vergrößerung der Wechselwirkungsrate.

Mit dieser Motivation wurde 1952 von D. A. Glaser und L. Alvarez die *Blasen-kammer* entwickelt, die mit einer fast siedenden Flüssigkeit gefüllt wird, so dass bei plötzlicher Expansion kurzzeitig Überhitzung eintritt. An den Ionen bilden sich Dampfbläschen, die durch geeignete Beleuchtung sichtbar gemacht werden. Die Füllflüssigkeit ist zumeist Wasserstoff, der zunächst bei einem Druck von 5 bis 6 bar schwach unterkühlt und völlig blasenfrei ist. Dann wird innerhalb kürzester Zeit der Druck auf die Hälfte reduziert. Für wenige Millisekunden, nämlich bis das Sieden einsetzt, ist die Blasen-kammer für Teilchen, die durch dünne Metallfenster eintreten, empfindlich.

Die Blasen-kammer leistete erhebliche Dienste bei der Entdeckung neuer Teilchen, unter ihnen zum Beispiel *seltsame* Teilchen wie das Kaon. Nebel- und Blasen-kammern werden in einem Magnetfeld betrieben, denn aus der Krümmung der entstehenden Bahnen kann auf den Impuls der auslösenden Teilchen geschlossen werden. Die Nebelkammer wurde vollständig durch die Blasen-kammer verdrängt, die später oft sogar mit supraleitenden Magneten versehen wurden.

1959 bauten S. Fukui und S. Miyamoto die erste Funkenkammer. Durchquert ein Teilchen die mit Edelgas gefüllte Kammer, ionisiert es entlang seiner Flugbahn die Gasatome. Im Falle eines kompletten Durchquerens der Kammer wird über eine Steuerungslogik eine Hochspannung auf leitende Platten in der Kammer gelegt, um diese Ladungsträger zu beschleunigen. Es kommt zu einer lawinenartigen Vermehrung. Sind genügend Ladungsträger entstanden, so gibt es eine leitende Verbindung und die Spannung bricht in Form eines Funken zusammen. Viele parallele Platten sorgen dafür, dass längs der Strahlungsspur eine Funkensequenz zu beobachten ist. Die Teilchenbahnen sind zwar nicht so scharf gezeichnet wie in der Blasen-kammer, dafür hat die Funkenkammer gegenüber der Blasen-kammer zwei Vorzüge. Zunächst einmal verfügt man bei der Verwendung von Platten aus Schwermetall über ein sehr massives Target, was zum Nachweis von Teilchen mit besonders geringer Wechselwirkung mit Materie von Vorteil ist. So ermöglichte erst die Einführung der Funkenkammer viele Experimente zur Neutrino-Physik, denn durch das dichte Targetmaterial wurde die Wahrscheinlichkeit einer Neutrinoreaktion groß genug, um Experimente durchzuführen. Deutlich wird die Bedeutung dieser Möglichkeit, wenn man bedenkt, dass die Wahrscheinlichkeit für Wechselwirkung mit Materie für Neutrinos so gering ist, dass sie die Erde größtenteils ungestört durchqueren können.

Der zweite Vorteil von Funkenkammern ist, dass sie bei einer relativ kurzen Totzeit elektronisch *getriggert*, d. h. gesteuert ausgelöst werden können. Mit unterschiedlichen Techniken lässt sich zunächst feststellen, ob ein Ereignis von Interesse ist. Nur in diesem Fall wird die Spannung an die Platten der Funkenkammer angelegt.

Die Funkenkammer kommt heute nicht mehr in der Forschung zur Teilchenphysik zum Einsatz. Sie dient meistens als Demonstrationsobjekt und erlaubt es, natürliche Radioaktivität für jeden sichtbar darzustellen.

1. Theoretische Grundlagen der Teilchenphysik

Die Beschäftigung der Menschheit mit der Frage nach der inneren Struktur der Materie hat eine lange Geschichte, die sich mit Unterbrechungen zweieinhalb Jahrtausende bis zu den griechischen Philosophen der Antike zurückverfolgen lässt. Viele der ehemals als elementar bezeichneten Teilchen können heute nicht mehr als Grundbausteine der Materie angesehen werden. Ebenso wie der ursprünglich von Demokrit für den Urbaustein eingeführte Begriff *Atom* (griech.: *ἄτομος* = nicht teilbar) seinem Sinn nicht mehr gerecht wurde – es zeigte sich, dass Substrukturen existieren – wandelte sich mit neuen Erkenntnissen auch immer das Bild vom elementaren Teilchen.

Das Ziel der Teilchenphysik ist es, die Eigenschaften und Bestandteile der Materie auf möglichst grundlegende Prinzipien zurückzuführen. Dabei sind in den letzten hundert Jahren immer wieder entscheidende Fortschritte gelungen.

So gelangte man zum Beispiel Anfang des 20. Jahrhunderts durch die Experimente von E. Rutherford, H. Geiger und E. Marsden zum modernen Bild des Atoms, wonach es aus einem positiv geladenen Kern besteht, in dem seine Masse konzentriert ist und der von einer Elektronenwolke umgeben ist [Rut11]. Der Kern lässt sich wiederum in noch kleinere Teilchen zerlegen, Protonen und Neutronen. In den dreißiger Jahren waren nur wenige Elementarteilchen bekannt. Dies waren neben Elektron, Proton und Neutron nur noch das Photon. Das Neutrino war aus theoretischen Erwägungen postuliert.

Die Untersuchung kosmischer Strahlung und Beobachtung der Reaktionen an immer leistungsfähigeren Beschleunigern eröffnete jedoch bald den Blick auf eine enorme Vielfalt von Teilchen, die der Begriff „Teilchen-Zoo“ verdeutlicht. Insbesondere zeigte die tief inelastische Elektron-Proton-Streuung, dass auch das Proton keineswegs elementar ist, sondern eine Substruktur besitzt. Diese und andere Entdeckungen führten zu einem relativ einfachen Bild der elementaren Materiebausteine und deren Wechselwirkungen, dem so genannten *Standardmodell* [PRS99].

1.1 Der Aufbau der Materie

1.1.1 Quarks und Leptonen

Im Rahmen der Teilchenphysik werden die kleinsten Strukturen untersucht. Heute wird der Aufbau der Materie im Standardmodell aus einfachen, punktförmigen ($< 10^{-18}$ m) Konstituenten beschrieben, den nach heutigem Verständnis elementaren Teilchen.

Diese Konstituenten werden in zwei Klassen eingeteilt, die *Leptonen*, deren bekanntester Vertreter das Elektron ist und die *Quarks*, deren wichtigste Vertreter das u- (up) und das d- (down) Quark sind (siehe Tab. 1.1).

Fermionen	Familie			elektrische Ladung	Farbe	Spin
	1	2	3			
Leptonen	ν_e	ν_μ	ν_τ	0	–	$1/2$
	e	μ	τ	–1		
Quarks	u	c	t	$+2/3$	r,b,g	$1/2$
	d	s	b	$-1/3$		

Tabelle 1.1: Eigenschaften der fundamentalen Fermionen [PRS99]. Zu jedem dieser Teilchen existiert ein entsprechendes Antiteilchen.

Diese fundamentalen Teilchen, Quarks und Leptonen, haben einen Eigendrehimpuls, auch *Spin* genannt, von $1/2$ und sind somit *Fermionen*. Die Quarks tragen die Ladung $+2/3$ bzw. $-1/3$ der Elementarladung, während die Leptonen ganzzahlige Ladungen besitzen. Zu allen diesen Fermionen gibt es die entsprechenden Antifermionen mit gleicher Masse, Zerfallszeit und Spin, aber entgegengesetzter elektrischer Ladung, Leptonen- (bzw. Baryonen-)zahl und entgegengesetztem magnetischem Moment.

Aus den sechs Quarksorten (*Flavors*) und ihren Antiquarks lassen sich die bekannten *Hadronen* aufbauen. Ein *Baryon* besteht im Rahmen des Standardmodells dabei meistens aus drei Quarks und ein *Meson* aus einem Quark-Antiquark-Paar. Neuere Experimente deuten auch auf mögliche 5-Quark-Zustände hin [Kub04].

Das *Pauli-Prinzip* besagt, dass zwei Fermionen nicht den gleichen Quantenzustand besetzen, also nicht in allen Quantenzahlen übereinstimmen. Im Quarkmodell sind auch Kombinationen drei gleicher Quarks möglich, die auf Grund des Pauli-Prinzips verboten sein müssten. Dass solche Teilchen trotzdem experimentell beobachtet wurden, hatte die Einführung einer weiteren Quantenzahl, der so genannten

Wechselwirkung	koppelt an	Austauschteilchen	Masse in GeV/c^2
stark	Farbe	8 Gluonen (g)	0
elektromagnetisch	elektrische Ladung	Photon (γ)	0
schwach	schwache Ladung	$W^\pm, Z^0$	$\approx 10^2$

Tabelle 1.2: Elementare Wechselwirkungen im Standardmodell der Teilchenphysik [PRS99].

Farbladung, zur Folge, damit das Pauli-Prinzip erhalten bleibt. Die Tatsache, dass keine freien Quarks beobachtet werden können, lässt sich in diesem Bild so beschreiben, dass nur farbneutrale Kombinationen der Quarks als freie Teilchen existieren (*Confinement*). Baryonen können als Kombination dreier Quarks mit den Farben *rot*, *gelb* und *blau* angesehen werden, während Mesonen aus einem Quark und dessen Antiquark, die eine Farbe und dessen Antifarbe tragen, zusammengesetzt sind. Die Theorie der stark wechselwirkenden Quarks ist die Quanten-Chromo-Dynamik (QCD), die in Anlehnung an die Quanten-Elektro-Dynamik (QED) zur Beschreibung elektrisch geladener Teilchen so genannt wird und ebenfalls als Eichtheorie formuliert ist.

Die elementaren Teilchen können durch den Austausch verschiedener fundamentaler *Bosonen* (Teilchen mit ganzzahligem Spin) wechselwirken. Diese so genannten *Austauschteilchen* agieren als Feldquanten der jeweiligen Wechselwirkung.

1.1.2 Wechselwirkungen

Neben der Gravitation gibt es drei elementare Wechselwirkungen, die in ihrer Struktur sehr ähnlich sind (siehe Tab. 1.2). Bei der elektromagnetischen Wechselwirkung ist es das Photon, das Kräfte zwischen geladenen Teilchen bewirkt. In der Theorie der starken Wechselwirkung sind es acht unterschiedliche Bosonen, die Gluonen, die die Kräfte zwischen den Farbladungen vermitteln. Die Austauschteilchen der schwachen Wechselwirkung sind das $W^\pm$ - und das Z^0 -Boson.

Die Gravitationskraft ist im Vergleich zu anderen Wechselwirkungen in der Teilchenphysik die schwächste. Die nächst stärkere ist die schwache Wechselwirkung. Sie ist für relativ langsame Prozesse, wie zum Beispiel β -Zerfälle von Kernen, verantwortlich.

Die elektromagnetische Wechselwirkung hat wie die Gravitation eine unendlich große Reichweite, wobei die Kraft, die das Feld einer Punktladung auf eine andere Ladung ausübt, mit dem Abstand abnimmt. Auf Grund der großen Reichweite ist

sie für die meisten Phänomene außerhalb des Kerns verantwortlich, so zum Beispiel für die Bindungszustände in Atomen und Molekülen.

Die starke Wechselwirkung hält die Quarks in den *Hadronen* zusammen. In einem Atomkern zum Beispiel bindet die starke Wechselwirkung die Quarks im Proton bzw. Neutron und die Restwechselwirkung der drei Quarks erzeugt die Kraft zwischen den Nukleonen. Das heißt, sie ist verantwortlich für deren Bindung im Kern. Ihre Stärke ist daran zu erkennen, dass selbst die Restwechselwirkung die Coulombabstoßung der Protonen überwindet.

Die schwache und die starke Wechselwirkung haben kurze Reichweiten; die der starken Wechselwirkung liegt mit ca. 10^{-15} m in der Größenordnung des Nukleonradius und die der schwachen ist mit ca. 10^{-18} m noch einmal drei Größenordnungen kleiner. Die typischen Zerfallszeiten, die für die starke Wechselwirkung $\sim 10^{-23}$ s, für die elektromagnetische Wechselwirkung $\sim 10^{-20}$ s und für die schwache Wechselwirkung $\sim 10^{-8}$ s betragen [Roe96], können ebenfalls als Maß für die Stärke der jeweiligen Wechselwirkung angesehen werden.

1.1.3 Symmetrien und Erhaltungssätze

Ein bestimmter Typ von Wechselwirkung gehorcht im Allgemeinen vielen verschiedenen Erhaltungssätzen, so dass die mathematische Beschreibung dieser Wechselwirkung unter mehreren Operationen invariant sein muss. Bei allen Wechselwirkungen bleiben Energie, Impuls, Drehimpuls, Ladung, Farbe, Baryonenzahl und die drei Leptonenzahlen erhalten. Eine Symmetrie unter dem Operator X bedeutet, dass das System gleiches Verhalten nach der Anwendung von X zeigt. In der Teilchenphysik sind die Erhaltung der *Parität* und die der *Ladungskonjugation* sehr wichtige Symmetrien. Durch den Paritätsoperator P wird die Raumspiegelung am Koordinatenursprung erzeugt. In der Elementarteilchenphysik bedeutet dies, dass ein Prozess im Spiegel genauso ablaufen würde (gleiche Wahrscheinlichkeiten, Produkte, Edukte). Der Eigenwert dieses Operators (wenn er existiert) ist ± 1 , und dieser Wert wird die Parität des Systems genannt.

Die Ladungskonjugation, die mathematisch durch den Operator C beschrieben wird, bedeutet eine Änderung der Vorzeichen aller ladungsartigen Quantenzahlen ($q, B, L, S, \dots$) und transformiert im Allgemeinen ein Teilchen in das zugehörige Antiteilchen. Eine Symmetrie unter Ladungskonjugation bedeutet gleiche Zerfallszeiten, Aufspaltungsverhältnisse sowie gleiche Teilcheneigenschaften für die ladungskonjugierten Zustände.

Die Parität P und Ladungskonjugation C bleiben bei der starken und der elektromagnetischen Wechselwirkung erhalten, nicht jedoch bei der schwachen Wechsel-

wirkung. Der Nachweis der Paritätsverletzung bei der schwachen Wechselwirkung gelang erstmals S. Wu et al. 1957 durch Untersuchungen des β -Zerfalls von ^{60}Co [Wu57].

Bei der Paritätsoperation ändert auch die so genannte *Helizität* ihr Vorzeichen. Die Helizität H ist als normiertes Produkt aus Spin $\mathbf{s}$ und Impuls $\mathbf{p}$

$$H = \frac{\mathbf{s} \cdot \mathbf{p}}{|\mathbf{s}| \cdot |\mathbf{p}|} \quad (1.1)$$

definiert. Durch die Normierung kann die Helizität H nur die Eigenwerte $+1$ oder -1 annehmen. Teilchen mit Spin in Bewegungsrichtung haben $\langle H \rangle = +1$, Teilchen mit Spin entgegen der Bewegungsrichtung haben $\langle H \rangle = -1$. Teilchen mit $\langle H \rangle = +1$ werden als rechtshändig bezeichnet und solche mit $\langle H \rangle = -1$ als linkshändig.

Der geladene Strom der schwachen Wechselwirkung ist maximal paritätsverletzend, denn er koppelt nur an linkshändige Fermionen und rechtshändige Antifermionen; der neutrale schwache Strom ist teilweise paritätsverletzend, da er in unterschiedlicher Stärke an links- und rechtshändige Fermionen und Antifermionen koppelt. Da nur linkshändige Neutrinos und rechtshändige Antineutrinos existieren, kann weder $P|\nu_L\rangle \rightarrow |\nu_R\rangle$ noch $C|\nu_L\rangle \rightarrow |\bar{\nu}_L\rangle$ existieren. Auf Grund der Verletzung von P und C ist es sinnvoll, die kombinierte CP -Parität zu betrachten, denn sie führt Neutrinos in die physikalisch relevanten Antineutrinos über. Jedoch ist im neutralen Kaon-System auch die Kombination von CP nicht erhalten. Wird als weitere Komponente noch die Zeitumkehr T hinzu genommen, wobei Zeitinvarianz, also Symmetrie unter T , heißt, dass die Zeitrichtung eines Prozesses nicht erkennbar ist, dann erhält man eine Größe CPT . Das so genannte CPT -Theorem besagt, dass jede relativistische lokale Feldtheorie unter CPT erhalten sein muss.

Nur der geladene Strom der schwachen Wechselwirkung wandelt Quarks in andere Quarks (Quarks mit anderem Flavor) und Leptonen in andere Leptonen um. Die Quantenzahlen, die den Quark-Flavor angeben (dritte Komponente des Isospins, Strangeness, Charm etc.) bleiben daher bei den übrigen Wechselwirkungen erhalten. Der Betrag des Isospin bleibt bei der starken Wechselwirkung erhalten.

1.2 Wechselwirkung von Strahlung mit Materie

Durchqueren geladene Teilchen Materie, so geben sie durch Stöße mit dem Medium Energie ab. Am häufigsten handelt es sich hier um Wechselwirkungen mit den Hüllenelektronen der Atome, die zur Anregung oder Ionisation führen. Der mittlere Energieverlust eines geladenen Teilchens wird durch die *Bethe-Bloch-Formel*

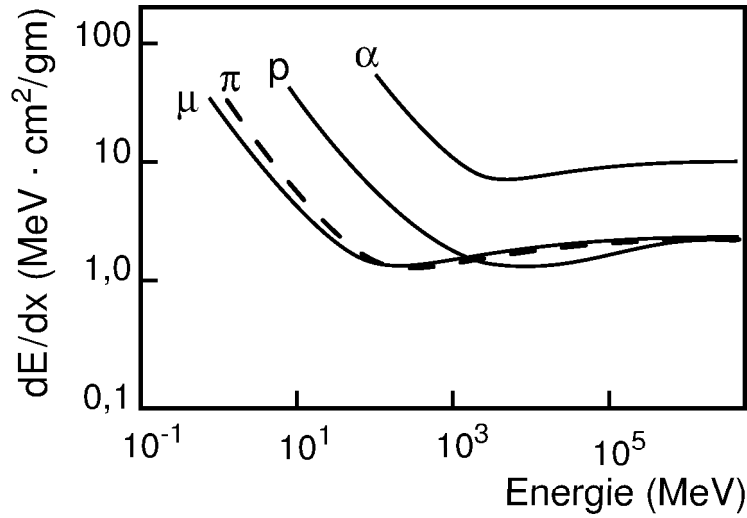


Abbildung 1.1: Mittlerer Energieverlust pro durchquerter Weglänge als Funktion der Energie für einige Teilchen nach [Leo87].

beschrieben:

$$-\frac{dE}{dx} = \frac{4\pi N_0 z^2 e^4}{mv^2} \frac{Z}{A} \left[\ln \frac{2mv^2}{I \cdot (1 - \beta^2)} - \beta^2 \right], \quad (1.2)$$

wobei m die Elektronenmasse und z die Ladung des Teilchens bezeichnet. Die Größe β ist die auf die Lichtgeschwindigkeit c normierte Geschwindigkeit v des Teilchens ($\beta = v/c$). N_0 ist die Loschmidt-Avogadro-Zahl, Z und A ist die Ordnungszahl bzw. Massenzahl der Atome des Materials. Die Größe I bezeichnet ein effektives Ionisationspotenzial, gemittelt über alle Elektronen. Bis auf diese Größe sind alle Parameter in der Bethe-Bloch-Gleichung gut bestimmt. Das Potenzial ist eine atomphysikalische Größe, die im Prinzip quantenmechanisch berechnet werden kann. Wegen der Vielzahl der möglichen atomaren Anregungen ist dies im Allgemeinen sehr schwierig und in der Praxis wird daher auf empirische Werte zurückgegriffen. Näherungsweise gilt: $I = 16 \text{ eV} \cdot Z^{0.9}$ für Kernladungszahlen $Z > 1$ [PRS99].

Die Größe x steht für die Dicke der durchquerten Schicht in Einheiten von g/cm^2 . Wird der Energieverlust pro dx mit $dx = \rho y$ angegeben, wobei ρ die Dichte und y die Dicke des Absorbers in Richtung der Flugbahn des Teilchens ist, hat dies den Vorteil, dass der Energieverlust für viele unterschiedliche Materialien kaum variiert.

Eine wichtige Folgerung aus der Gleichung von Bethe und Bloch ist die Energieabhängigkeit des Energieverlustes pro Weglänge dE/dx , der unabhängig von der Masse M des Teilchens ist. In Abb. 1.1 ist der Energieverlust pro Weglänge

als Funktion der kinetischen Energie für verschiedene Teilchen aufgetragen. Für nicht-relativistische Geschwindigkeiten fällt dE/dx wie $1/v^2$ ab, da die Zeit für die Wechselwirkung zwischen der Punktladung des Teilchens und der Ladungsverteilung des Absorbers abnimmt, und erreicht ein Minimum bei einer Geschwindigkeit von $v \approx 0,96 c$. Der minimale Wert von dE/dx ist für alle Teilchen gleicher Ladung ähnlich. Teilchen aus diesem Energiebereich bezeichnet man auf Grund des minimalen Energieübertrags als *minimal ionisierend*. Für Myonen zum Beispiel liegt der Energieverlust durch Ionisation pro Weglänge dx im Ionisationsminimum und auch bei höheren Teilchenenergien bei etwa $2 \text{ MeV}/(\text{g}\cdot\text{cm}^{-2})$.

Mit weiter zunehmender Geschwindigkeit spielen relativistische Effekte eine Rolle. Die Längenkontraktion führt zu einer scheinbaren Erhöhung der Ladungsdichte, mit der das Teilchen wechselwirkt. Folglich steigt dE/dx wieder an, jedoch lediglich logarithmisch mit dem Dilatationsfaktor $\gamma = E/Mc^2 = (1 - \beta^2)^{-1/2}$.

Generell kann die Bethe-Bloch-Formel nur für eine genäherte Berechnung des Abbremsvermögens herangezogen werden. Für die exakte Berechnung werden insbesondere bei hohen und niedrigen Projektilgeschwindigkeiten materialabhängige Korrekturterme eingeführt.

Der größte Teil des Energieverlustes führt zur Bildung von Ion-Elektron-Paaren im Material. Zwei Stufen dieses Prozesses werden unterschieden. In der ersten erzeugt das einfallende Teilchen in atomaren Stößen primäre Ionen. Die dabei herausschlagenden Elektronen haben eine Verteilung in der Energie E' von der Form $dE'/(E')^2$ [Per90]. Die so entstandenen Elektronen höherer Energie können in der zweiten Stufe auf ihrem Weg durch das Material ebenfalls Ionen erzeugen (Sekundärionisation). So ist die Gesamtzahl von Ionen drei- bis viermal größer als die der primären Ionen; sie ist proportional zum Energieverlust des einfallenden Teilchens im Material.

Die Bethe-Bloch-Formel (1.2) gibt nur den mittleren Wert des Energieverlustes dE in der Schichtdicke dx an, der reale Wert ist statistisch verteilt, wobei die Fluktuationen von der relativ kleinen Anzahl *harter* Primärstöße mit großem E' bestimmt wird. Aus diesem Grund ist die Verteilung um den Mittelwert eine *Landau-Verteilung*, also asymmetrisch mit einem Ausläufer zu hohen Werten. Trotzdem ist es möglich, den mittleren Energieverlust auf wenige Prozent genau zu bestimmen, indem zum Beispiel in vielen aufeinander folgenden dicken Gasschichten der Energieverlust gemessen wird und die größten Messwerte nicht berücksichtigt werden. So wird der Dilatationsfaktor γ eines einfallenden Teilchens aus seiner Ionisationswirkung bestimmt. Mit zusätzlicher Information des Impulses kann auf diese Weise die Masse des Teilchens abgeschätzt und somit eine Teilchenidentifikation vorgenommen werden.

2. Kosmische Strahlung

Die Erde wird permanent von hochenergetischen Teilchen aus dem Weltall getroffen. Dieses Phänomen wurde 1912 von V. F. Hess entdeckt [Hes36]. Zuvor war die natürliche Radioaktivität zwar bereits bekannt, aber man ging davon aus, dass die an der Erdoberfläche gemessene ionisierende Strahlung nur von radioaktiven Nukliden in der Erdkruste verursacht wird. In einem Heißluftballon stieg V. F. Hess bis auf Höhen von 5000 m auf, wo er mit Hilfe mehrerer Elektrometer Intensitätsmessungen der ionisierenden Strahlung in unterschiedlichen Höhen durchführte und entdeckte, dass die Intensität der Strahlung mit zunehmender Höhe ansteigt. Anfangs wurde die Sonne als mögliche Quelle angenommen, doch eine totale Sonnenfinsternis 1927 lieferte keine nennenswerten Unterschiede zu Messungen außerhalb der Sonnenfinsternis. Somit schied die Sonne als dominierender Faktor aus und der Ursprungsort der kosmischen Strahlung musste außerhalb unseres Sonnensystems vermutet werden.

V. F. Hess folgerte, dass uns diese Strahlung aus den Tiefen des Universums erreicht und gab ihr den Namen *kosmische Strahlung*. Für seine Untersuchungen wurde er 1936 mit dem Nobelpreis in Physik ausgezeichnet.

2.1 Die primäre kosmische Strahlung

Die Rate der primären kosmischen Strahlung, die die Erdatmosphäre erreicht, liegt insgesamt bei rund 1000 Teilchen pro m^2 und Sekunde und teilt sich auf in Atomkerne und eine Reihe von Elementarteilchen, wobei die Atomkerne mit 98 % den Hauptbestandteil bilden. Sie teilen sich auf in 87 % Wasserstoffkerne (Protonen), 12 % Heliumkerne und 1 % Kerne schwererer Elemente.

Die Erde wird von etwa zehn Teilchen pro m^2 und Minute getroffen, die eine Energie von einem TeV besitzen. Schon bei der 100fachen Energie verringert sich diese Häufigkeit auf ungefähr fünf Teilchen pro m^2 und Tag. Im Bereich oberhalb von etwa 10 PeV wird nur noch ein Teilchen pro m^2 und Jahr gemessen. Der Fluss bei Energien über 10^{20} eV liegt lediglich bei ca. einem Teilchen pro km^2 und Jahrhundert.

Am oberen Rand der Atmosphäre ist die Energieverteilung der Strahlung nicht mehr die gleiche wie im Weltraum. Sie wird durch den Entstehungsprozess bestimmt und ändert sich im erdmagnetischen Feld, denn Teilchen kleiner Energie werden abgelenkt und können unter Umständen die Erdatmosphäre nicht mehr erreichen.

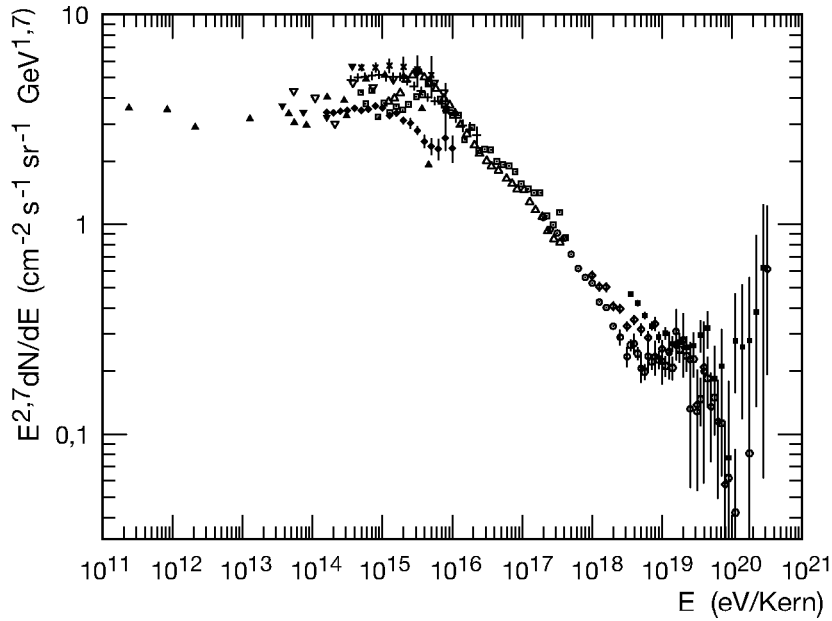


Abbildung 2.1: Differenzielles Energiespektrum der primären kosmischen Strahlung multipliziert mit $E^{2,7}$. Die unterschiedlichen Datenpunkte stammen aus verschiedenen Messungen nach [PDG00].

Die untere Grenze der Energie der auf die Erdatmosphäre treffenden Teilchen hängt dabei außer von der Teilchensorte noch von der geografischen Breite und der Richtung der Teilchen ab. In mittleren Breiten liegt sie für Protonen in der Größenordnung $3 \cdot 10^9$ eV. Für höhere Energien entspricht die Energieverteilung jener im interstellaren Raum.

Die Häufigkeit der Teilchen, die die Erde erreichen, nimmt mit steigender Teilchenenergie stark ab. Die kosmische Strahlung überdeckt ein Energieintervall von einigen MeV bis hin zu Energien über 10^{20} eV. Unterhalb von einigen GeV überlagert sich das Spektrum der kosmischen Strahlung aus der Galaxie mit der Strahlung aus unserer Sonne. Oberhalb von 10 GeV bis hin zu den höchsten bis heute gemessenen Energien lässt sich das Spektrum mit einem Potenzgesetz der Form

$$\frac{dN}{dE} = \text{const } E^{-\gamma} \quad (2.1)$$

beschreiben, dessen Parameter allerdings abschnittsweise angepasst werden müssen.

So gibt es einen Knick im Spektrum bei einigen PeV, der als *Knie* bezeichnet wird. Dabei ändert sich der Exponent (spektraler Index) γ von etwa 2,7 unterhalb des Knies auf ungefähr 3,1 darüber. Eine zweite Modifikation am Potenzspektrum ist bei einer Energie von etwa $5 \cdot 10^{18}$ eV notwendig. Dort scheint sich ein Abflachen des Flussspektrums anzudeuten. Dieser Bereich wird in Analogie zum Knie der *Knöchel*

genannt. Abbildung 2.1 zeigt das differentielle Energiespektrum in Abhängigkeit von der Teilchenenergie, welches zur Verdeutlichung der Strukturen mit $E^{2,7}$ multipliziert ist.

Sowohl über die Herkunft des Knies als auch über den Ursprung des Knöchels gibt es keine gesicherten Erkenntnisse. Eine von zahlreichen Spekulationen über das Zustandekommen des Knies geht von verschiedenen Ursprungsorten der kosmischen Strahlung der Bereiche unter- und oberhalb des Knies aus. So könnte die Quelle der Strahlung von galaktischer zu extragalaktischer Herkunft wechseln. Es wird sowohl eine Änderung des Beschleunigungsmechanismus der kosmischen Strahlung, als auch eine drastische Änderung der Elementzusammensetzung an der Quelle in Betracht gezogen.

So wie das Zustandekommen des Knies bis heute noch nicht richtig verstanden ist, ist auch das Phänomen des Knöchels Gegenstand aktueller Forschung. Eine mögliche Ursache könnte der so genannte *Greisen-Zatsepin-Kuz'min Effekt* sein [Mül03]. Protonen mit einer Energie über 10^{19} eV können mit der 2,7 K-Hintergrundstrahlung wechselwirken. Ihre Energie reicht aus, um bei Stoßprozessen mit den Photonen Pionen oder Elektron-Positron-Paare zu bilden. Dadurch verliert die kosmische Strahlung einen Großteil ihrer Energie. Da die mittlere freie Weglänge von Protonen dieser hohen Energien somit beschränkt ist, müsste es zu einer Anreicherung von Teilchen unterhalb dieser GZK-Energie kommen. Der Fluss von Teilchen oberhalb von 10^{19} eV sollte demnach deutlich verringert sein. Die Teilchen sollten auf Grund der geringen Reichweite aus unserer unmittelbaren Umgebung stammen. Bis jetzt konnte aber noch kein Objekt in der Nachbarschaft unserer Galaxis mit einem Mechanismus beobachtet werden, der Teilchen auf solch hohe Energien beschleunigen könnte [Gre66, Lon92]. Klärung verspricht man sich unter anderem von dem in Argentinien neu aufgebauten Luftschauerexperiment AUGER. Auf einer Gesamtfläche von 3000 km^2 sollen noch in diesem Jahr Wasser-Čerenkov-Detektoren und Teleskope für Fluoreszenzlicht diese höchstenergetische kosmische Strahlung untersuchen.

Bei der Suche des Ursprungs kosmischer Strahlung bzw. des Mechanismus, der Teilchen auf Energien bis über 10^{20} eV beschleunigen kann, ist es von Bedeutung, dass die meisten geladenen Teilchen von interstellaren Magnetfeldern beeinflusst werden und komplizierte Bahnen beschreiben. Daher ist die Verteilung der Einschläge in der Erdatmosphäre isotrop und es gehen Informationen über den Entstehungsort verloren. Teilchen der kosmischen Strahlung mit Energien größer 10^{19} eV werden nicht mehr von diesen Magnetfeldern beeinflusst und bewegen sich geradlinig. Zur Zeit ist die Anzahl der registrierten Teilchen jedoch noch zu gering, um statistisch signifikante Aussagen zu ermöglichen.

Messungen haben ergeben, dass der Fluss der kosmischen Strahlung zeitlich konstant ist [Lon92]. Demnach muss gewährleistet sein, dass mögliche Quellen die Energiedichte der Strahlung aufrecht erhalten. Als Quellen der kosmischen Strahlung kommen Pulsare, Neutronensterne, Schwarze Löcher und Supernovae in Frage. Die Existenz der hohen und höchsten Energien ist zwar experimentelle Tatsache, aber ihre Herkunft ist noch ungeklärt.

Informationen über die Quellen der kosmischen Strahlung bieten die Elementzusammensetzung, das Energiespektrum sowie die Untersuchung der elektromagnetischen Komponente, welche direkte Richtungsinformationen liefern kann (zum Beispiel durch *Gamma Ray Bursts*, GRB [Fis96]).

2.2 Wechselwirkung mit der Erdatmosphäre

Kosmische Strahlung trifft auf die Erdatmosphäre und erzeugt dort Teilchenschauer. Die hochenergetischen Primärteilchen dringen in die Atmosphäre ein und erzeugen über die starke Wechselwirkung mit Atomkernen (Stickstoff, Sauerstoff,...) niederenergetische Sekundärteilchen, die ihrerseits mit den Atomen der Atmosphäre wechselwirken, wodurch *hadronische Schauer* entstehen (siehe Abb. 2.2).

So werden in diesen Nukleon-Nukleon-Wechselwirkungen unter anderem Pionen gebildet, die auf Grund ihrer kurzen Lebensdauer den Erdboden nicht erreichen. Neutrale Pionen zerfallen mit einer mittleren Lebensdauer von $\tau_{\pi^0} = 8,4 \cdot 10^{-17}$ s [PDG00] sehr schnell über die elektromagnetische Wechselwirkung. Geladene Pionen zerfallen über die schwache Wechselwirkung, ihre mittlere Lebensdauer beträgt $\tau_{\pi^\pm} = 2,6 \cdot 10^{-8}$ s [PDG00]. Das π^- ist das leichteste Hadron mit elektrischer Ladung und kann nur in Leptonen zerfallen:

$$\begin{aligned}\pi^- &\rightarrow \mu^- + \bar{\nu}_\mu, \\ \pi^- &\rightarrow e^- + \bar{\nu}_e.\end{aligned}$$

Die entsprechenden Zerfälle des π^+ ergeben sich durch Ladungskonjugation. Der zweite Prozess ist gegenüber dem ersten im Verhältnis 1:8000 unterdrückt, obwohl der zweite Kanal vom Phasenraum bevorzugt wird. Dieses Verhalten ist auf die Paritätsverletzung der schwachen Wechselwirkung zurückzuführen (siehe Kap. 1.1.3).

Das Pion zerfällt in nur zwei Teilchen, die im Schwerpunktsystem in entgegengesetzte Richtungen emittiert werden. Da es Spin 0 hat, muss sich auch das entstandene Lepton-Neutrino-Paar in einem Zustand mit Gesamtspin $S = 0$ befinden. Die Helizität des Antineutrinos liegt fest; es ist rechtshändig, hat also immer positive Helizität: $\langle H \rangle = +1$. Auf Grund der Drehimpulserhaltung muss der Spin

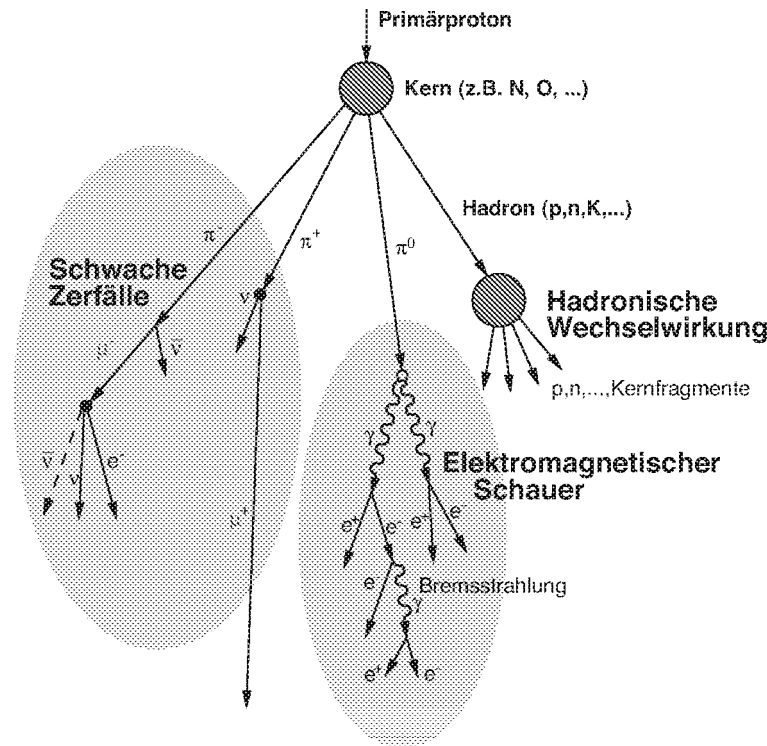


Abbildung 2.2: Wechselwirkung der kosmischen Strahlung mit der Erdatmosphäre.

des Leptons entgegengesetzt orientiert sein, also in Bewegungsrichtung des Leptons, welches somit auch positive Helizität haben muss. Dies ist aber die „falsche“ Helizität für das Lepton, um an das W-Boson, das Austauschteilchen der schwachen Wechselwirkung, zu koppeln. Für masselose Fermionen koppelt der geladene schwache Strom ja nur an die linkshändigen Fermionen (siehe Kap. 1.1.3). Auf Grund ihrer endlichen Masse haben Elektronen und Myonen mit Spin in Bewegungsrichtung aber auch eine linkshändige Komponente, die proportional zu $1 - \beta$ ist. An diese Komponente koppelt das W-Boson. Die Elektronenmasse ist kleiner als die des Myons ($m_e = 0,511 \text{ MeV}/c^2$, $m_\mu = 106 \text{ MeV}/c^2$). So beträgt beim Pionenzerfall $1 - \beta_e = 2,6 \cdot 10^{-5}$ und $1 - \beta_\mu = 0,72$.

Die linkshändige Komponente des Elektrons ist also viel kleiner als die des Myons und der Zerfall entsprechend stark unterdrückt. Wird das Verhältnis dieser beiden Werte gebildet und zusätzlich bedacht, dass der zur Verfügung stehende Phasenraum die Unterdrückung zugunsten des Zerfalls in ein Elektron um den Faktor 3,5 mindert, erhält man das Verhältnis von ca. 1:8000 [PRS99].

2.3 Das Myon

Durch Wechselwirkung von kosmischen Teilchen mit der Atmosphäre entsteht eine Vielzahl von Sekundärteilchen. Der die Erdoberfläche erreichende Anteil der geladenen Teilchen setzt sich auf Meereshöhe aus ca. 79 % Myonen, 20 % Elektronen und etwa 1 % anderer Teilchen zusammen [Gru85]. Das Elektron und das Myon sind die leichtesten elektrisch geladenen Teilchen. Daher und auf Grund der Ladungserhaltung ist das Elektron stabil und beim Zerfall des Myons muss ein Elektron entstehen. Das Myon zerfällt schwach durch Austausch eines W-Bosons gemäß:

$$\mu^- \rightarrow e^- + \bar{\nu}_e + \nu_\mu.$$

Auch hier ergeben sich die entsprechenden Zerfälle des μ^+ durch Ladungskonjugation. In seltenen Fällen wird dabei zusätzlich ein Photon oder ein e^+e^- -Paar erzeugt. Myonen haben eine mittlere Lebensdauer von $\tau_\mu = 2,197 \cdot 10^{-6}$ s [PDG00], weshalb aus der klassischen Mechanik folgt, dass sie bei einer gemessenen Geschwindigkeit von $v \approx 0,99983 \cdot c$ (c = Lichtgeschwindigkeit) im Mittel nur eine Strecke von ca. 660 m zurücklegen.

Da sich Myonen nahezu mit Lichtgeschwindigkeit bewegen, muss hier aber relativistisch gerechnet werden. Die Relativitätstheorie besagt, dass auf Grund der endlichen Lichtgeschwindigkeit die Zeit mit zunehmender Geschwindigkeit für einen Beobachter immer langsamer verläuft. Deshalb erfahren die hier behandelten Myonen eine Zeitdilatation, die ihre Lebensdauer für einen Beobachter auf der Erde verlängert. Mit der oben genannten Geschwindigkeit v wird hier zur Anschauung die zurückgelegte Strecke bestimmt:

Klassisch:

$$v(\text{Myon}) = 0,99983 \cdot c$$

$$\Rightarrow \text{Strecke, die es im Mittel zurücklegt: } s = v \cdot \tau = 660 \text{ m}$$

Relativistisch:

Für Beobachter:

Mittlere Lebensdauer eines Myons mit der Geschwindigkeit $v(\text{Myon}) = 0,99983 \cdot c$:

$$\tau' = \frac{\tau}{\sqrt{1 - \frac{v^2}{c^2}}} = 1,1 \cdot 10^{-4} \text{ s}$$

$$\Rightarrow \text{Strecke, die es im Mittel zurücklegt: } s = v \cdot \tau' = \mathbf{32,9 \text{ km}}$$

Die Reichweite der Myonen hängt auf Grund der Wechselwirkungen mit der Atmosphäre auch von ihrer Energie ab (siehe Kap. 1.2). Die meisten Myonen werden

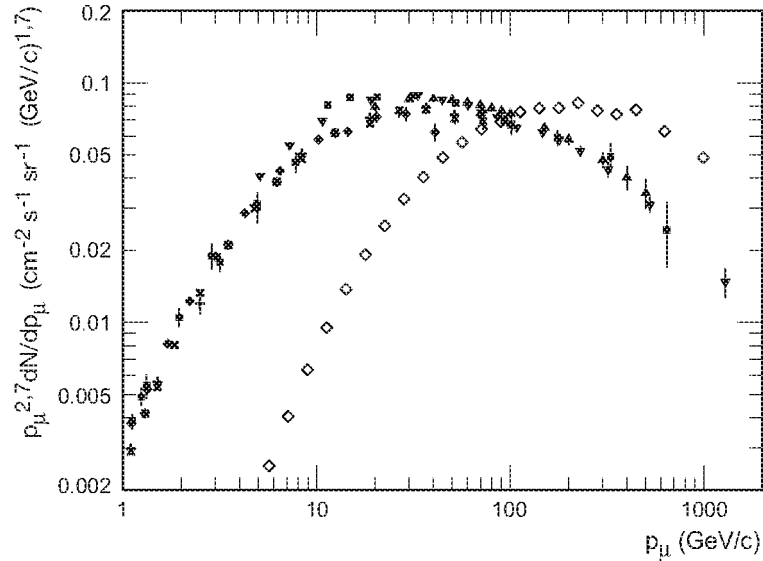


Abbildung 2.3: Myonenspektrum bei $\theta = 0^\circ$ (schwarze Symbole) und $\theta = 75^\circ$ (weiße Rauten). Die unterschiedlichen Datenpunkte stammen aus verschiedenen Messungen nach [PDG00].

relativ hoch in der Atmosphäre, ungefähr in 15 km Höhe, produziert und verlieren durch Ionisation eine Energie von etwa 2 GeV, bevor sie den Boden erreichen. Die mittlere Energie der Myonen am Boden beträgt ungefähr 4 GeV [PDG00]. Der gesamte Myonenfluss auf Meereshöhe liegt bei ca. 200/m²s [Gru85], wobei die Energie der Myonen stark vom Einfallswinkel abhängig ist. Wegen der unterschiedlichen durchquerten Strecke der Atmosphäre bei geneigten Winkeln variiert der totale Myonenfluss mit dem Zenitwinkel θ wie

$$I_\mu(\theta) \sim \cos^2 \theta, \quad (2.2)$$

wobei θ von der Vertikalen gerechnet wird.

Dieser Zusammenhang gilt nicht mehr für sehr hochenergetische Myonen. Falls der Energieverlust der Myonen in der Atmosphäre vernachlässigbar im Vergleich zur Myonenenergie wird, ändert sich auch deren Zenitwinkelabhängigkeit. Für Myonenenergien größer 5 TeV wird

$$I_\mu(\theta) \sim 1/\cos \theta. \quad (2.3)$$

Die höchstenergetischen Myonen fallen fast horizontal ein. Das liegt daran, dass die Pionen, die die Hauptquelle für die Myonen bilden, aus schräg in die Atmosphäre einfallender Primärstrahlung einen längeren Weg (proportional zu $1/\cos \theta$) durch

weniger dichte Atmosphäre zurücklegen. Daher zerfallen sie mit größerer Wahrscheinlichkeit, bevor sie an einer Wechselwirkung teilnehmen, wobei quasi ihr Gesamtimpuls auf das Myon (und das Myon-Neutrino) übertragen wird. Pionen aus vertikalen Schauern treten in den hohen Dichtegradienten ein und geben eher Energie durch Wechselwirkung ab bevor sie zerfallen. So verteilen sie ihre Energie auf viele niederenergetische Teilchen.

Abb. 2.3 zeigt das Energiespektrum für Myonen auf Meeresniveau für zwei Winkel. Bei großen Einfallswinkeln zerfallen niederenergetische Myonen, bevor sie die Erdoberfläche erreichen und hochenergetische Pionen zerfallen bevor sie wechselwirken. Daher steigt die mittlere Myonenenergie.

Insgesamt erreichen die Erdoberfläche mehr positiv als negativ geladene Myonen. Das Verhältnis reflektiert den Überschuss von π^+ zu π^- durch die Fragmentation von Protonen in Vorwärtsrichtung sowie den Überschuss von Protonen zu Neutronen im primären Spektrum. Das Verhältnis beträgt zwischen 1,1 und 1,4 für Energien von 1 GeV bis 100 GeV. Für Myonenenergien unter 1 GeV ändert sich das Verhältnis abhängig von geomagnetischen Effekten systematisch mit dem Ort [PDG00].

3. Das Grundprinzip der Funkenkammer

Die Funkenkammer besteht aus einem mit Edelgas (75% Neon und 25% Helium bei Normaldruck) gefüllten Gefäß, in dem sich eine Reihe von parallelen, Kupfer beschichteten Platten befindet, die als Elektroden fungieren. Die Platten sind alternierend über Kondensatoren mit einer Hochspannungsquelle (HV) oder direkt mit dem Erdpotenzial verbunden (siehe Abb. 3.1).

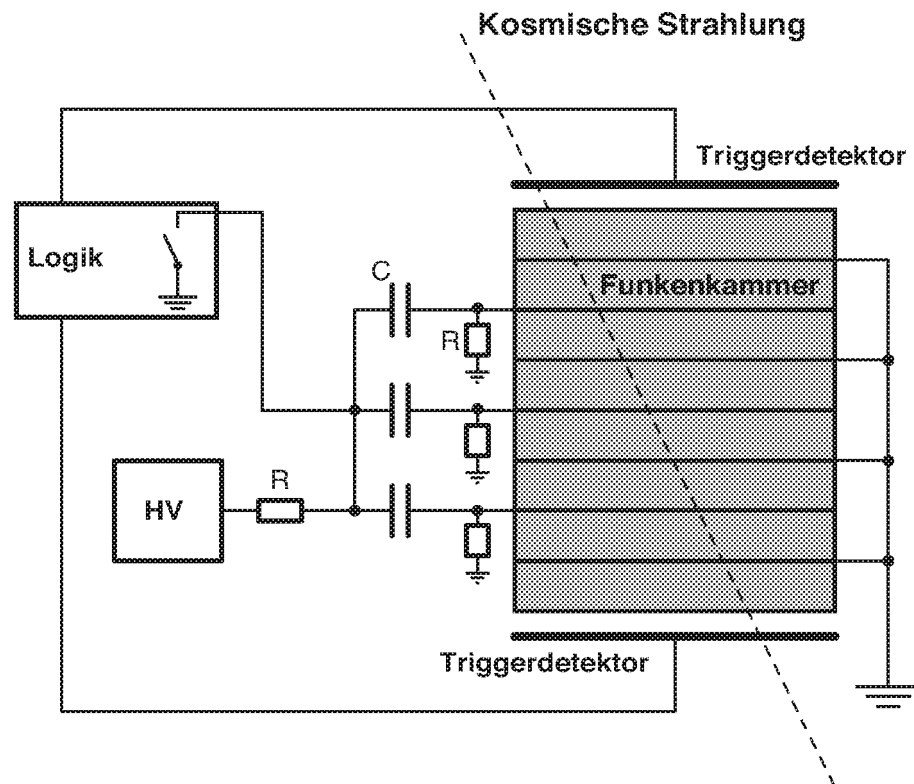


Abbildung 3.1: Prinzipaufbau der Funkenkammer.

Nach dem Durchgang eines ionisierenden Teilchens wird über eine Schaltlogik ein Hochspannungsimpuls an jede zweite Platte gelegt. Zur Detektion der Teilchen sind zwei Triggerdetektoren an der Kammer angebracht (siehe Kap. 4), die im Fall des Teilchendurchgangs durch die Kammer ein Signal an die Schaltlogik geben, welche in Kombination mit den Kondensatoren als Schalter für die Hochspannung dient (siehe Kap. 5).

Liegt die Hochspannung an den Platten, werden die entstandenen Ladungsträger durch das elektrische Feld räumlich getrennt. Die Ionen driften zur Kathode und die Elektronen werden in Richtung der positiv geladenen Anode beschleunigt. Die Energie der Elektronen reicht aus, um weitere Gasatome zu ionisieren. So kommt es zu einer lawinenartigen Vermehrung von Ladungsträgern und am Ort der Primäriionisation bildet sich innerhalb von wenigen Nanosekunden eine Entladung aus (siehe Kap. 3.1). Der Entladungskanal verläuft im Allgemeinen parallel zum elektrischen Feld und so ist es möglich, die komplette Spur eines ionisierenden Teilchens durch die Kammer anhand der entstandenen Funken nachzuvollziehen (siehe Abb. 3.2).

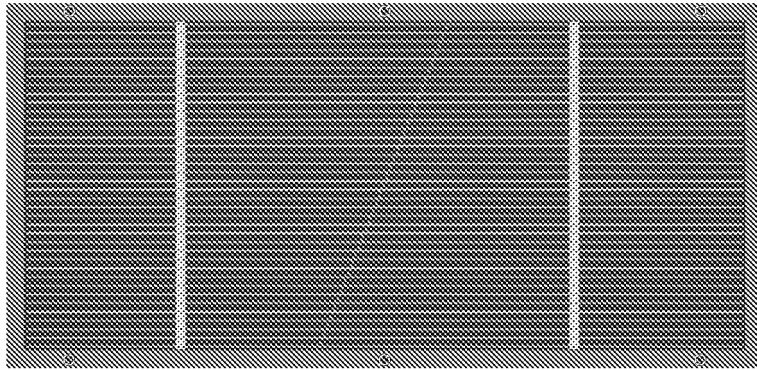


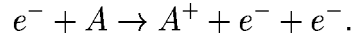
Abbildung 3.2: Zeichnung der Funkenkammer mit einer Abfolge von Funken.

Das Ansprechvermögen der Funkenkammer (die Wahrscheinlichkeit für Funkenbildung nach einem Teilchendurchgang) hängt, wegen der Diffusion der primären Ladungsträger, von der zeitlichen Verzögerung zwischen dem Teilchendurchgang und dem Zeitpunkt des Hochspannungsimpulses ab. Da beim Aufbau dieser Kammer auf ein elektrisches Reinigungsfeld verzichtet wurde, das permanent anliegen würde, wird die Diffusionsbewegung der Ladungsträger nicht durch eine feldbedingte Driftbewegung überlagert; sie entfernen sich also nur langsam vom Ort der Primäriionisation und das Anlegen der Hochspannung etwa $1 \mu\text{s}$ nach dem Teilchendurchgang ist vollkommen ausreichend (siehe Kap. 3.2).

3.1 Ionisierung eines Gases

Geladene schnelle Teilchen werden hauptsächlich auf Grund ihrer Ionisationswirkung unterschieden. Gasatome oder -moleküle lassen sich ionisieren, indem thermische, optische, elektrische oder kinetische Energie zugeführt wird. Werden zum Beispiel Elektronen auf genügend hohe Energien beschleunigt, so können sie beim Stoß mit

Atomen oder Molekülen ein Elektron aus der Elektronenhülle herausschlagen und dadurch ein Elektron-Ion-Paar erzeugen:



Dies ist der Hauptmechanismus zur Erzeugung von Ladungsträgern in Gasentladungen. Um durch Stoßionisation neue Ladungsträger zu erzeugen, müssen die Elektronen im beschleunigenden Feld $\mathbf{E}$ während der freien Weglänge Λ zwischen zwei Stößen mindestens eine Energie aufnehmen, die ausreicht, um das gestoßene Gasatom mit der Ionisationsenergie W_{ion} zu ionisieren (siehe dazu auch [Dem95]). Bei einem elektrischen Feld $\mathbf{E}$ in x -Richtung ist ihre Energieaufnahme $e \cdot E \cdot \Lambda_x$ und die Bedingung für Stoßionisation auf der Strecke Λ_x ist:

$$e \cdot E \cdot \Lambda_x \geq W_{\text{ion}}. \quad (3.1)$$

N Elektronen, die im Feld $\mathbf{E}$ in x -Richtung beschleunigt werden, erzeugen dann entlang der Strecke dx

$$dN = \gamma N dx \quad (3.2)$$

neue Ladungsträgerpaare und damit dN zusätzliche Elektronen, die nach entsprechender Beschleunigung wieder stoßionisieren können. Der Faktor

$$\gamma = \frac{dN/N}{dx} \quad (3.3)$$

gibt die Anzahl der Sekundärelektronen an, die ein Primärelektron im Mittel pro Weglängeneinheit in x -Richtung erzeugt und wird als *Ionisierungsvermögen* bezeichnet. Es ist vom Verhältnis E/p , also von Feldstärke E und Druck p sowie von der Ionisationsenergie W_{ion} im Gas abhängig.

Durch Integration der Gleichung (3.2) für den Elektronenstrom ergibt sich die nach der Strecke x angewachsene Zahl von Elektronen pro Zeiteinheit zu

$$N = N_0 e^{\gamma x}, \quad (3.4)$$

wobei $N_0 = N(x=0)$ die Anzahl der vorhandenen Elektronen bei $x=0$ ist.

Das Ionisationsvermögen γ wird nach der *Townsend-Theorie* auch *erster Townsend-Koeffizient* genannt [Far65].

Untersuchungen u.a. mit Hilfe fotografischer Methoden ergeben folgende Vorstellung von der Entstehung eines Funkens (siehe Abb. 3.3): Eine Lawine von Ladungsträgern, die sich ausgehend von der primären Ionisation in Feldrichtung ausbreitet,

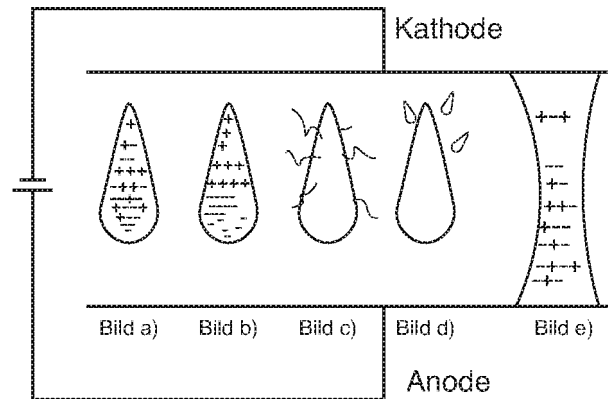


Abbildung 3.3: Die Entstehung eines Streamers.

bildet wegen der verschiedenen Beweglichkeit der Elektronen und Ionen eine *Tropfenform* (siehe Abb. 3.3a). Die sich wesentlich schneller bewegenden Elektronen nehmen den Kopf ein und hinterlassen eine Wolke positiv geladener Ionen (siehe Abb. 3.3b). Auf diese Weise entsteht zwischen Ionen und Elektronen ein zusätzliches elektrisches Feld, welches das Äußere in einem kleinen Raumbereich abschwächt. Dagegen wird das elektrische Feld zwischen der Anode und den Elektronen sowie der Kathode und den Ionen verstärkt.

Hat sich die Verstärkung N/N_0 der Ladungsträger auf über 10^6 erhöht, ist das Raumladungsfeld zwischen Ionen und Elektronen groß genug, um das äußere Feld soweit zu kompensieren, dass im feldfreien Raum Rekombination stattfinden kann. Kommt es dazu, werden UV-Photonen emittiert (siehe Abb. 3.3c). Diese Photonen ionisieren in der Nähe der ursprünglichen Entladung Gasatome. Die freigesetzten Elektronen driften in Richtung der positiven Ionen, wobei es im Bereich hoher Feldstärken an der Spitze der Ionenwolke zu Ladungsvervielfachungen kommt; es entstehen weitere sekundäre Lawinen (siehe Abb. 3.3d). Weitere Photonen entstehen durch Rekombination, so dass sich die Entladung in Richtung der Kathode ausbreitet. Erreicht die Verstärkung der Ladungsträger einen kritischen Wert von 10^8 (*Raether-Bedingung* [Rae63]), ist das Raumladungsfeld von gleicher Größe wie das äußere Feld und es folgt nach Gleichung (3.4):

$$\frac{N}{N_0} = \exp(\gamma x_{\text{RB}}) = \exp(20) \approx 10^8. \quad (3.5)$$

Sekundäre und primäre Lawinen bilden zusammen *Streamer*, die sich in Richtung der beiden Elektroden ausdehnen (siehe Abb. 3.3e). Die Streamer wachsen entlang der Bahnspur zu einem Plasmakanal zusammen und wenn sie die Elektroden erreichen, entsteht eine leitende Verbindung, die als Funken sichtbar wird [Gru96].

Der hier beschriebene Prozess wird *Streamer-Prozess* genannt und benötigt zur Ausbildung eines Funkens eine Zeit von etwa 10^{-8} s [All69].

Bevorzugt werden Gase, die schon bei niedriger Spannung einen hohen Townsend-Koeffizienten haben und somit die Raether-Bedingung (3.5) erfüllen, da in solchen Gasen nach einer wesentlich kürzeren Driftstrecke x_{RB} das Verhältnis der Anzahl der Ladungsträger N zur Anzahl der primären Ladungsträger N_0 den Wert von 10^8 erreicht. In diesem Fall folgt aus Gleichung (3.5):

$$x_{RB} \cdot \gamma \approx 20. \quad (3.6)$$

Abbildung 3.4 zeigt den von Gasdruck und -art abhängigen ersten Townsend-Koeffizienten γ pro Druckeinheit p als Funktion vom Verhältnis E/p für verschiedene Gase und für Luft.

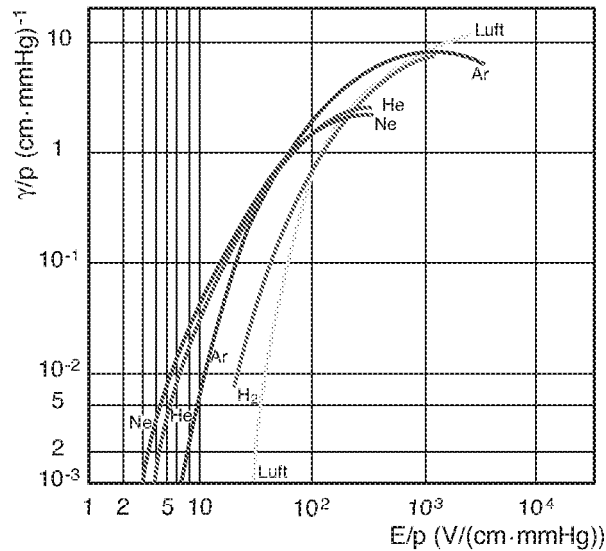


Abbildung 3.4: Erster Townsend-Koeffizient γ auf den Druck normiert als Funktion des Verhältnisses E/p für verschiedene Gase nach [Ric74].

Die als Elektroden fungierenden Platten in der Funkenkammer haben zueinander einen Abstand d von 0,8 cm. Bei einer Driftstrecke von $x = d$ ist nach Gleichung (3.6) bei einem Townsend-Koeffizienten von $\gamma = 25 \text{ cm}^{-1}$ die Raether-Bedingung erfüllt. Da die Funkenkammer bei Normaldruck von 760 Torr (760 mmHg) betrieben werden soll, ergibt sich auf den Druck normiert ein Wert für γ/p von $0,03 \text{ (cm mmHg)}^{-1}$. Abbildung 3.4 ist zu entnehmen, dass bei den Arbeitsbedingungen der Funkenkammer für Neon und Helium im Vergleich zu anderen Gasen das kleinere elektrische Feld nötig ist. Für $\gamma/p = 0,03$ nimmt E/p in einer Neon- und Helium-Umgebung Werte von etwa $4\text{--}6 \text{ V/(cm mmHg)}$ an, für Luft hingegen Werte von über 25 V/(cm mmHg) .

Damit lassen sich die elektrischen Felder abschätzen, die bei Normaldruck zur Funkenbildung erforderlich sind: etwa 5 kV/cm für Neon und Helium bzw. 19 kV/cm für Luft.

Um bei einem Plattenabstand von 0,8 cm ein elektrisches Feld von 5 kV/cm zu erzeugen, muss eine Spannung von 4 kV angelegt werden. Zur Erzeugung eines Feldes von 19 kV/cm ist schon eine Spannung von über 15 kV erforderlich. Die für die Funkenkammer benutzte Hochspannungsquelle liefert jedoch lediglich Spannungen bis 8 kV. So ist für eine mit Luft gespülte Kammer ein Betrieb unter diesen Bedingungen nicht möglich, da zur Funkenentladung wesentlich höhere elektrische Felder erforderlich sind. Darüber hinaus haben die Entladungen in einer Neon-Atmosphäre eine auffallend rote Färbung, die für die optische Beobachtung sehr geeignet ist. So wird für die Funkenkammer eine Mischung aus Neon und Helium verwendet.

Jede Verunreinigung mit Luft setzt die zur Funkenentladung notwendige Spannung herauf; bei der Inbetriebnahme der Funkenkammer ergab sich eine optimale Hochspannung von 4600 V. Dieser Wert deckt sich mit den Erwartungen aus der Theorie.

3.2 Diffusion in feldfreiem Gas

Die mittlere thermische Energie eines Gasmoleküls mit drei Freiheitsgraden beträgt $\epsilon_T = 3 \cdot \frac{1}{2} kT$. Dieses folgt aus dem *Gleichverteilungssatz*, welcher besagt, dass auf jeden Freiheitsgrad im thermischen Gleichgewicht die gleiche mittlere Energie entfällt. Bei Normalbedingungen ($T = 273$ K) beträgt $\epsilon_T \approx 0,035$ eV. Die Verteilung der kinetischen Energie ϵ bei der Temperatur T ist die *Maxwell-Boltzmann-Verteilung*:

$$F(\epsilon) = c\sqrt{\epsilon} \cdot e^{-\epsilon/kT}. \quad (3.7)$$

Eine zum Zeitpunkt $t = 0$ punktförmig am Ursprung lokalisierte Ladungsverteilung diffundiert durch Vielfachstreuung in den umgebenden Raum, wobei sich eine Gaußverteilung um den Ursprung bildet. Wird die differentielle Dichteverteilung dN/N von Ladungsträgern in Abhängigkeit vom Ort x und zum Zeitpunkt t betrachtet, kann somit der Diffusionskoeffizient D definiert werden:

$$\frac{dN}{N} = \frac{1}{\sqrt{4\pi Dt}} e^{-(x^2/4Dt)} dx. \quad (3.8)$$

Der Diffusionskoeffizient ist umso größer, je größer die mittlere thermische Geschwindigkeit u ist, die mit der *Idealen Gaskonstante* R und der *Molaren Masse* M als

$$u = \sqrt{\frac{3RT}{M}} \quad (3.9)$$

definiert ist [Mor96]. Sie ist im Allgemeinen die Geschwindigkeit eines Atoms oder Moleküls, dessen kinetische Energie dem Mittelwert der kinetischen Energien aller Atome oder Moleküle im Gas entspricht. In Helium beträgt demnach die mittlere Geschwindigkeit der Ladungsträger $1,3 \cdot 10^5$ cm/s und in Neon $0,58 \cdot 10^5$ cm/s. Genähert entfernen sich die Ionen also in einer Zeit von $1 \mu\text{s}$ eine Strecke von ca. 1 mm vom Ort der Primärionisation. Der Wert von $1 \mu\text{s}$ ist ein typischer Wert für die Verzögerungsdauer zwischen Teilchendurchgang und Anwesenheit des elektrischen Feldes, da sich das Anlegen der Hochspannung in dieser Zeit mit den zur Verfügung stehenden Bauteilen gut realisieren lässt.

Durch Rekombination kann es ebenfalls zur Abnahme der Dichte der Ladungsträger kommen. Positive Ionen rekombinieren mit Elektronen, bevor sie das elektrische Feld in entgegengesetzte Richtungen beschleunigt.

Einen Beleg dafür, dass die Ladungsdichte $1 \mu\text{s}$ nach der Primärionisation auf jeden Fall noch hoch genug ist, liefert eine Messung von J. G. Rutherglen, die in Abb. 3.5 dargestellt ist.

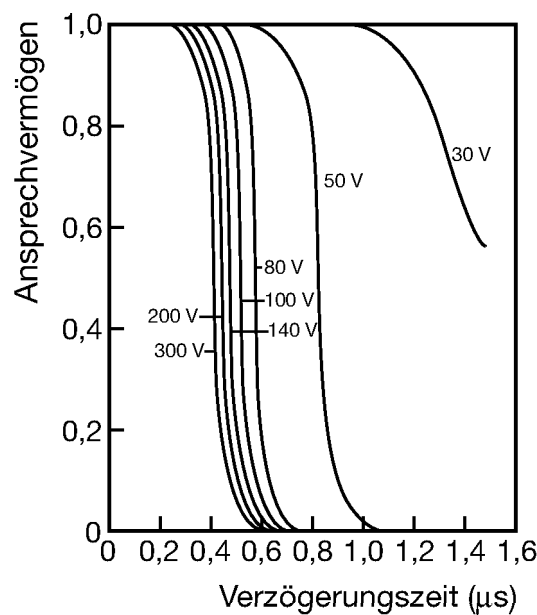


Abbildung 3.5: Ansprechvermögen (d. h. Wahrscheinlichkeit für Funkenbildung) einer Funkenkammer als Funktion der Verzögerungszeit für verschiedene Werte der Reinigunsspannung [Rut64].

Dieses Diagramm zeigt das Ansprechvermögen einer Funkenkammer als Funktion der Verzögerungszeit für verschiedene Werte der Reinigunsspannung. Dabei bedeutet ein Ansprechvermögen von eins, dass in der Funkenkammer bei jedem Triggersignal die Bahnspur in Form von Funkenentladungen zu beobachten ist. Ein Verringern dieser Reinigunsspannung vergrößert die mögliche Verzögerungszeit, da

die durch die Spannung zustande kommenden Driftbewegungen der Ladungsträger abnehmen. Ohne Reinigungsspannung liegt man mit einer Verzögerungszeit von $1\ \mu\text{s}$ auf jeden Fall in einem Bereich akzeptablen Ansprechvermögens.

3.3 Der Gasaustausch mit der Umgebung

Die Erzeugung der Funken in der Funkenkammer hängt stark von der Reinheit des Gasgemisches ab, da Verunreinigungen lokal die Zündspannung erhöhen können (insbesondere Luft, siehe Kap. 3.1). Aus diesem Grund muss die Kammer vor der Inbetriebnahme mehrere Tage lang mit einer Rate von mehreren l/h mit dem Neon-Helium-Gemisch gespült werden. Es dringt über eine Öffnung in die Kammer ein und strömt aus einer weiteren Öffnung auf der anderen Seite der Rückwand wieder aus. Mehrere Male wird dabei das gesamte Volumen ersetzt, bevor Bahnspuren in Form von Funken beobachtbar sind. Wird die Kammer vor der Inbetriebnahme nicht ausreichend gespült, sind bei gleichen Parametern (Hochspannung, Druck, Plattenabstand) keine Funkenentladungen zu erkennen. Auf Grund der anfangs herrschenden Unreinheit der Gasmischung wäre als Voraussetzung für Funkenentladungen eine höhere Spannung an den Platten notwendig (siehe Kap. 3.1). Mit längerem Spülen verschwinden die Verunreinigungen im Gasgemisch und die notwendige Hochspannung für Funkenbildung sinkt.

An dieser Stelle sei erwähnt, dass sich bei sehr hohen Spannungen nach einem Zündvorgang häufig Entladungen unabhängig von den Bahnspuren ausbilden. Diese Entladungen finden oft an denselben Stellen statt, welche zum Beispiel kleine Spitzen im Material der Platten sein können. Auch die Rate dieser unerwünschten Entladungen sinkt mit längerem Spülen der Kammer (erfahrungsgemäß nach ca. einer Woche bei einem Fluss von 2 l/h).

Nach dem Spülvorgang wird die Kammer permanent mit einem niedrigeren Fluss von ungefähr 0,5 l/h mit Gas versorgt. Dabei spielt auch nach ausreichendem Spülen der Kammer die Gasundurchlässigkeit des Kammergefäßes eine große Rolle, um den Gasaustausch mit der Umgebung zu verhindern.

In Gemischen von Gasen, die nicht miteinander reagieren, setzt sich der Gesamtdruck p aus den Partialdrücken der einzelnen Gaskomponenten zusammen, welche für Diffusionsbewegungen entscheidend sind.

Der Partialdruck einer Komponente A entspricht dem Druck, den diese Komponente ausüben würde, wenn sie als einziges Gas in gleicher Menge im gleichen Volumen anwesend wäre. Der Gesamtdruck ist gleich der Summe der Partialdrücke aller Komponenten $p = p(A) + p(B) + p(C) + \dots$. Nach der kinetischen Gastheorie haben die Moleküle des Gases A die gleiche mittlere kinetische Energie wie die

des anderen Gases B, da beide die gleiche Temperatur haben und sie sich nach Voraussetzung gegenseitig nicht anziehen. Das Vermischen von zwei oder mehreren Gasen ändert nichts an der mittleren kinetischen Energie einer Komponente des Gases. So können zum Beispiel Gaskomponenten der Luft trotz gleichen Gesamtdrucks wie das Gasgemisch in der Kammer höhere Partialdrücke haben als die Gase in der Kammer, deshalb durch Lecks in die Kammer hineindiffundieren und somit das Gasgemisch verunreinigen. Der Partialdruck des Gases A ergibt sich aus seinem Stoffmengenanteil:

$$p(A) = \frac{n(A)}{n(A) + n(B)} \cdot p. \quad (3.10)$$

Zur Berechnung der Partialdrücke der Gase in der Kammer können diese bei Zimmertemperatur und Atmosphärendruck zur Vereinfachung als *ideale Gase* ($pV = nRT$) betrachtet werden. Das Volumen der Kammer beträgt 64 l und es herrscht dort näherungsweise Atmosphärendruck.

$$n(\text{He}) = \frac{pV}{RT} = \frac{101,325 \text{ kPa} \cdot 16 \text{ l}}{8,31 \text{ kPa} \cdot \text{l} \cdot \frac{1}{\text{mol}} \cdot \frac{1}{\text{K}} \cdot 293,15 \text{ K}} = 0,665 \text{ mol} \quad (3.11)$$

$$n(\text{Ne}) = \frac{pV}{RT} = \frac{101,325 \text{ kPa} \cdot 48 \text{ l}}{8,31 \text{ kPa} \cdot \text{l} \cdot \frac{1}{\text{mol}} \cdot \frac{1}{\text{K}} \cdot 293,15 \text{ K}} = 1,99 \text{ mol} \quad (3.12)$$

Daraus folgt für die Partialdrücke der Gase in der Kammer

$$p(\text{He}) = \frac{n(\text{He})}{n(\text{He}) + n(\text{Ne})} \cdot p = 25,4 \text{ kPa}, \quad (3.13)$$

$$p(\text{Ne}) = \frac{n(\text{Ne})}{n(\text{Ne}) + n(\text{He})} \cdot p = 75,9 \text{ kPa}. \quad (3.14)$$

Werden diese Partialdrücke mit denen in der Luft verglichen

$$\begin{aligned} p(\text{N}_2) &= 79,1 \text{ kPa}, \\ p(\text{O}_2) &= 21,2 \text{ kPa}, \\ p(\text{Ar}) &= 0,9 \text{ kPa etc.}, \end{aligned}$$

so fällt auf, dass zum Beispiel Stickstoff einen recht hohen Partialdruck hat und, auf Grund des Bestrebens nach Konzentrationsausgleich durch Lecks in der Kammer, hineindiffundieren kann. Bei der Inbetriebnahme der Funkenkammer fiel deutlich auf, dass die Bahnspuren in Form von Funken erst zu beobachten waren, nachdem die Kammer auf Gasundurchlässigkeit überprüft und vorhandene Lecks abgedichtet wurden.

3.4 Die Gasversorgung

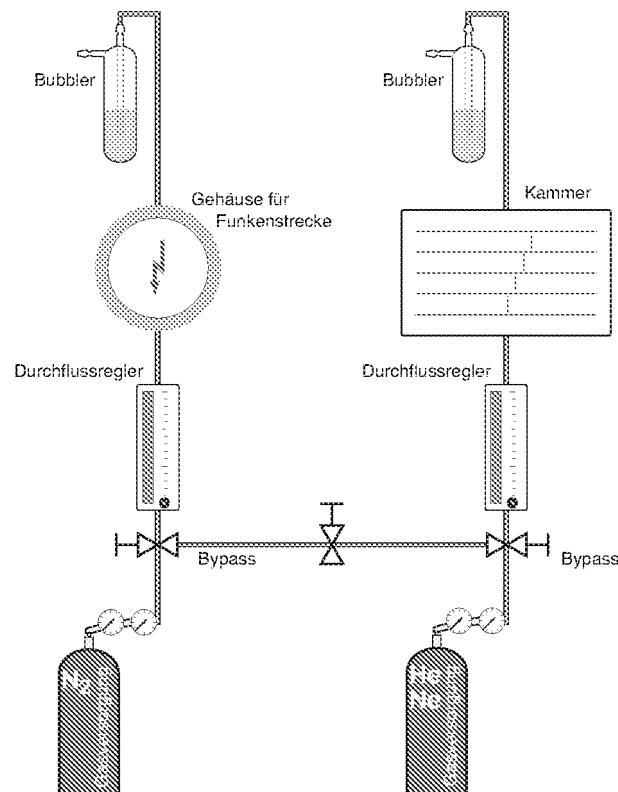


Abbildung 3.6: Schematische Darstellung der Gasversorgung.

Das Gassystem umfasst eine Flasche des Neon-Helium-Gemisches und eine Flasche Stickstoff jeweils mit einem Leervolumen von 10 l. Das Gas für die Funkenkammer wird im richtigen Mischungsverhältnis geliefert und in die Gasflasche umgefüllt. Mittels eines Druckreduzierers wird für das Neon-Helium-Gemisch ein Arbeitsdruck von 0,4 bar über Normaldruck eingestellt, mit dem es durch einen Durchflussregler in die Kammer fließt. Mit Stickstoff wird, ebenfalls durch einen Durchflussregler, ein Gehäuse für die Funkenstrecke der Zündschaltung versorgt (siehe Kap. 5.1).

Zwischen den beiden Leitungen sitzt ein Ventil, mit dem auch die Zufuhr von Stickstoff in die Kammer möglich wird. Dazu wird die Versorgung mit Helium und Neon durch Umlegen eines Bypasses gestoppt und das Ventil zwischen den Leitungen geöffnet. Diese Funktion wird zum ersten Spülen der Kammer vor der Inbetriebnahme benötigt, um den Gehalt anderer Luftkomponenten wie Sauerstoff und Kohlenstoffdioxid schon vorher zu minimieren.

Die abführenden Leitungen der Kammer und des Gehäuses für die Funkenstrecke enden jeweils in einem *Bubbler*, einem mit Glycerin gefüllten Gefäß, in dem der Gasaustritt in Form von Bläschen zu beobachten ist. Die Bubbler sind nach den Durchflussreglern die letzte Kontrolle für den Gasfluss.

3.5 Messung der Leckrate

Um die Gasundurchlässigkeit der Kammer zu ermitteln, erfolgt eine Abschätzung der *Leckrate* durch einen einfachen Versuchsaufbau mit einem Wasserglas als Bubbler (siehe Abb. 3.7).

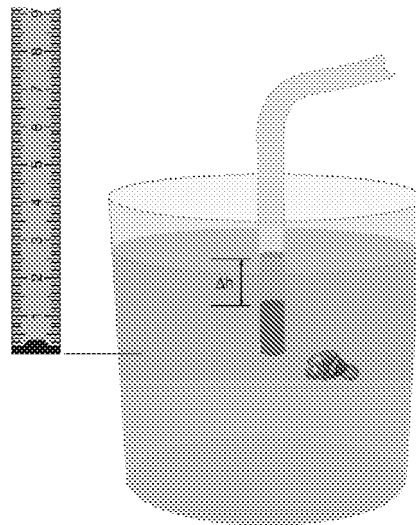


Abbildung 3.7: Versuchsaufbau zur Bestimmung der Leckrate.

Der Schlauch, durch den das Gas aus der Kammer strömt, wird mit seinem Ende ein paar Zentimeter tief in ein Wasserglas gegeben. Natürlich darf der Schlauch nicht zu tief in das Wasserglas getaucht werden, da sonst ein zu hoher Druck in der Kammer herrscht, der im schlimmsten Fall sogar noch weitere Lecks in den geklebten Kanten der Kammer verursacht. Bei offener Gaszufuhr treten Gasbläschen aus. Wird die Gaszufuhr gestoppt, steht die Gassäule im Schlauch bis zu dessen Ende still.

Der Atmosphärendruck, der außen auf der Wasseroberfläche lastet, drückt das Wasser im Schlauch hoch. Das Gasvolumen verkleinert sich, das Gas kann nur durch vorhandene Lecks der Kammer entweichen. Die Höhendifferenz zwischen dem Wasserspiegel innen und außerhalb des Schlauchs ist proportional zur Differenz des Drucks in der Kammer und dem Atmosphärendruck $\Delta p = \rho g \Delta h$. Daraus ergibt sich für eine Höhendifferenz der Wasserspiegel von 1 cm ein Druck von $\Delta p \approx 1 \text{ mbar}$.

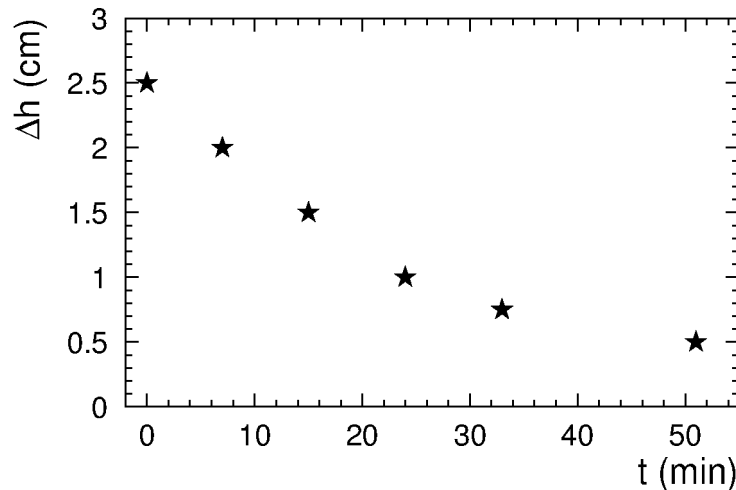


Abbildung 3.8: Höhendifferenz Δh der Wasserspiegel als Funktion der Zeit t .

Die Leckrate ist als Produkt aus dem Volumen der Kammer und dem Druckabfall pro Zeit definiert. Bei großem Überdruck ist der Druckabfall linear, bei kleineren Drücken fällt er immer flacher ab. In den nicht-linearen Bereich gelangt man unterhalb einer Druckdifferenz von ca. 1 mbar, wo die laminare Strömung durch vorhandene Lecks klein wird gegen die, von der Druckdifferenz unabhängigen, Diffusionsbewegungen (siehe Kap. 3.3).

So ist es auch nur sinnvoll, die Leckrate aus der Steigung des Grafen im linearen Bereich zu bestimmen:

$$\text{Maximale Leckrate} = \frac{p \cdot V}{t} = \frac{1,5 \text{ mbar} \cdot 64 \text{ l}}{24 \cdot 60 \text{ s}} = 0,07 \frac{\text{mbar} \cdot \text{l}}{\text{s}}. \quad (3.15)$$

4. Der Trigger

Viele Experimente werden *getriggert* betrieben, d. h. die Messung wird gezielt ausgelöst, wenn ein Ereignis von Interesse ist. Gerade für die Funkenkammer ist diese Möglichkeit von großer Bedeutung: Läge die Hochspannung permanent an den Platten der Funkenkammer, so käme es sehr häufig zu Funkenentladungen. So zum Beispiel auch, wenn Teilchen nur Sektionen der Kammer durchqueren. Darüber hinaus könnte die Spannung nicht permanent aufrecht erhalten werden. Aus diesen Gründen wird die Hochspannung nur bei Teilchendurchgängen durch die ganze Kammer an die Platten gelegt, d. h. die Kammer wird getriggert betrieben. Das bietet den Vorteil, dass ausschließlich interessante Ereignisse – Bahns Spuren durch die ganze Kammer – beobachtet werden und dass ein Hochspannungsimpuls während der *Totzeit* mit Hilfe eines *Timers* ausgeschlossen werden kann.

Die Realisierung erfolgt durch zwei Triggerdetektoren. Ober- und unterhalb der Kammer sind zwei *Szintillatoren* angebracht, die etwa die ganze Fläche der Kammer abdecken. Bei einem Kammerdurchflug passiert ein Teilchen beide Detektoren quasi gleichzeitig, da es sich näherungsweise mit Lichtgeschwindigkeit bewegt. Die geladenen Teilchen regen in den Szintillatoren Elektronen an und die absorbierte Energie wird in Form von UV-Licht wieder abgegeben, welches noch im Szintillator in sichtbares Licht umgewandelt wird. Diese Photonen werden über *Lichtleiter* auf das Fenster eines *Photomultipliers* (Fotoervielfacher) geleitet, der sie in einen elektrischen Spannungspuls umsetzt (siehe Abb. 4.1). Dieses *Triggersignal* ist der Auslöser für das Anlegen der Hochspannung. Im Folgenden werden die einzelnen Komponenten des Triggers detailliert beschrieben.

4.1 Der Szintillator und Lichtleiter

4.1.1 Die Funktionsweise von Szintillatoren

Szintillatoren spielen schon seit langer Zeit eine wichtige Rolle bei der Detektion einzelner geladener Teilchen. Für die Triggerdetektoren der Funkenkammer wird ein Plastikszintillator verwendet. Dieser gehört zu den organischen Szintillatoren, die sich durch sehr kurze Abklingzeiten des Szintillationslichts (im Bereich von Nanosekunden) auszeichnen. Ionisierende Teilchen regen beim Durchqueren spezielle Moleküle im Szintillator an. Beim Zerfall dieser angeregten Molekülniveaus entsteht UV-Strahlung. Da die Absorptionslänge von UV-Licht in den meisten organischen

Materialien sehr kurz ist (einige mm), gelingt die Extraktion eines Lichtsignals nur durch direkte Umwandlung des UV-Lichts in sichtbares Licht. Die Konversion dieses Lichts in blaues Licht ($\lambda \approx 430 \text{ nm}$) geschieht über Fluoreszenzstoffe, die so genannten Wellenlängenschieber. Diese organischen Moleküle werden ebenso wie die Szintillationsstoffe dem transparenten Trägermaterial beigemischt und so ausgewählt, dass die Absorption an die Wellenlänge des emittierenden Szintillators und die Emission an den Wellenlängenbereich der Empfindlichkeit der Fotokathode im Photomultiplier angepasst ist (siehe Kap. 4.2).

Die beiden aktiven Komponenten werden (wie bei dem hier verwendeten Szintillator) zum Beispiel mit Vinyltoluen, einem Monomer, zu einer polymerisierenden Substanz vermischt. Daraus wird dann Polyvinyltoluen, ein Feststoff, wobei die Polymerisierung in jede für praktische Anwendungen gewünschte Form erfolgen kann.

Das erzeugte Licht breitet sich innerhalb des Szintillators durch interne Totalreflexion aus. Um dieses Licht vom Ende des Szintillators auf das Fenster eines Photomultipliers zu leiten, werden geeignet geformte Lichtleiter verwendet.

4.1.2 Der Aufbau der Paddles

Die mit einem Photomultiplier als Triggerdetektor dienende Kombination aus Szintillator und Lichtleiter wird auf Grund ihrer Form *Paddle* genannt.

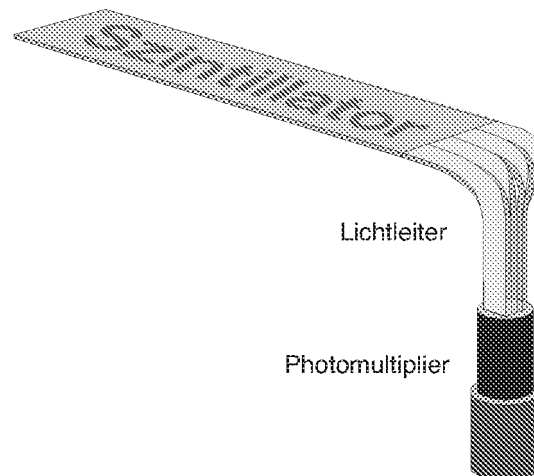


Abbildung 4.1: Skizze eines Triggerdetektors bestehend aus Paddle und Photomultiplier.

Die Größe der Plastiksintillatoren hängt von der gewünschten Detektionsfläche ab. Es ist sinnvoll, sie in möglichst gleichen Ausmaßen wie die Platten der Kammer zu wählen, denn so werden genau jene Teilchen detektiert, die auch die Kammer auf ganzem Weg durchdringen; die hier verwendeten Szintillatoren haben eine Fläche

von $72 \times 20 \text{ cm}^2$ und decken somit den Großteil der Plattenfläche von $80 \times 20 \text{ cm}^2$ ab.

Bei der Anfertigung der Lichtleiter muss beachtet werden, dass die rechteckige Szintillatorfläche vollständig auf das runde Fenster des Photomultipliers abgebildet wird. Darüber hinaus sollen die Paddles so geformt sein, dass sie mit den Photomultipliern platzsparend an der Kammer befestigt werden können. Diese Anforderungen erfüllen Lichtleiter aus gebogenen Plexiglasstreifen (siehe Abb. 4.1).

Auch hier erfolgt die Lichtleitung durch mehrere Totalreflexionen an den Oberflächen, die den Übergang vom Plexiglas zum umgebenden Medium bilden. Für den *Grenzwinkel* gilt nach Snellius das Brechungsgesetz:

$$\sin \alpha_g = \frac{n_2}{n_1}, \quad (4.1)$$

mit n_1 = Brechungsindex des Szintillators bzw. des Lichtleiters und n_2 = Brechungsindex des umgebenden Mediums.

Plexiglas hat einen Brechungsindex von $n_1 = 1,49$ und Luft $n_2 \approx 1$. Mit Gleichung (4.1) ergibt sich für den Lichtleiter ein Grenzwinkel α_g von 42° . Solange der Winkel des auf die Oberfläche treffenden Lichts größer als dieser Grenzwinkel α_g zur Normalen bleibt, ist der Übergang ins dünnere Medium, also der Austritt aus dem Lichtleiter nicht möglich; das Licht wird zu 100 % reflektiert. Das Licht sollte also unter einem Winkel größer als 42° zur Normalen auf die Fläche des Lichtleiters auftreffen, um größere Lichtverluste zu vermeiden. Deshalb muss der Krümmungsradius der Plexiglasstreifen groß gegen deren Dicke sein.

Für jeden Lichtleiter werden vier Plexiglasstreifen so zugeschnitten, dass ihre Gesamtfläche an einer Seite der Fläche der Szintillatorplatte entspricht ($A_1 = 1 \times 20 \text{ cm}^2$); an der anderen Seite muss die Fläche auf das kreisförmige Fenster des Photomultipliers mit dem Radius $r = 2,67 \text{ cm}$ abbildbar sein. Das Quadrat, das dort hinein passt, hat eine Fläche von $A_2 = 14,3 \text{ cm}^2$. Damit ändert sich bei gleichbleibender Dicke die Breite jedes Plexiglasstreifens von 5 cm auf 3,78 cm. Um diese so zu biegen, dass ihre Enden die richtige Form bilden, müssen sie auf 160°C erhitzt werden und dann in einer speziell angefertigten Halterung etwa eine Stunde lang fixiert werden. Die einzelnen Plexiglasstreifen werden nach dem Abkühlen mit Epoxidharz seitlich zusammengeklebt, damit sie nicht gegeneinander verrutschen können, was unebene Außenflächen und somit eine schlechtere Lichtkopplung zur Folge hätte. In Abb. 4.2 sind die in die richtige Form gebogenen Plexiglasstreifen zu sehen.

Die Lichtausbeute von Szintillator und Lichtleiter lässt sich deutlich steigern, wenn sie außen verspiegelt werden. Licht, das sonst am Materialrand auskoppeln

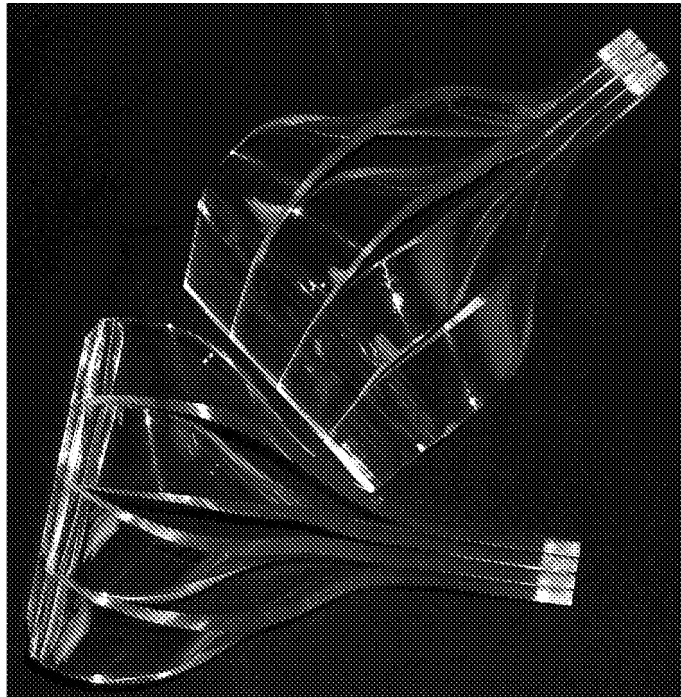


Abbildung 4.2: Gebogene Plexiglasstreifen für den Lichtleiter.

würde, wird dann reflektiert, so dass ein wesentlich größerer Teil des entstehenden Szintillationslichts genutzt wird. Deshalb werden die kompletten Paddles, inklusive der Szintillatoren, in Aluminium bedampfte Mylarfolie eingewickelt, die eine verspiegelte Fläche bietet und direkt am Paddle anliegt. Darüber werden die Paddles mit schwarzem Klebeband eingewickelt, was ein Eindringen von Licht verhindert und zusätzlich Stabilität bietet. Das ist wichtig, da die Übergänge zwischen Szintillatorplatten und Lichtleitern bzw. zwischen Lichtleitern und Photomultipliern nicht geklebt sind, damit die Bauteile ersetzt werden können. Sie sind mit optischem Fett aneinander gekoppelt, das den gleichen Brechungsindex wie der Plastiksintillator hat und somit die Anzahl weiterer Lichtbrechungen minimiert.

Die Szintillatorplatten werden mit den ankoppelnden Enden der Lichtleiter in eine zusätzlich stabilisierende Aluminiumverkleidung eingefasst, mit der sie später an der Kammer befestigt werden.

4.2 Der Photomultiplier

Mit einem Photomultiplier werden die Lichtimpulse aus dem Szintillator in elektrische Signale umgewandelt. Der Photomultiplier ist ein Sekundärelektronenvervielfacher (SEV) mit einer Fotokathode (siehe Abb. 4.3).

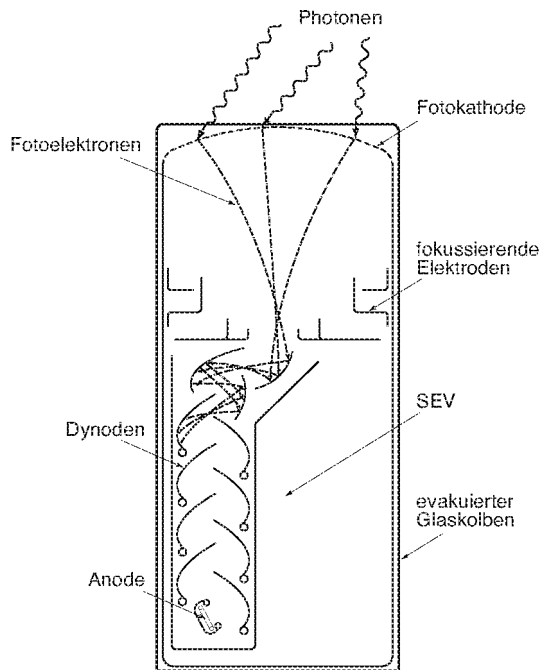


Abbildung 4.3: Typischer Aufbau eines Photomultipliers.

Die Photonen gelangen durch ein Fenster auf die in einer Hochvakuumröhre montierte Fotokathode, welche aus einer dünnen Schicht fotosensitiven Materials (hier eine Bialkali-Verbindung) besteht. Auftreffende Photonen können hier durch den fotoelektrischen Effekt Elektronen aus ihrer Bindung lösen.

Die Elektronen werden umso energiereicher, je höher die Frequenz ν des auslösenden Lichts ist. Da die Strahlungsenergie gequantelt ist, besitzen die Photonen eine Energie $h\nu$ (dabei ist h das *Plancksche Wirkungsquant*). Nach Abzug der Austrittsarbeit W_0 verbleibt dem Elektron noch die kinetische Energie von:

$$W = h\nu - W_0. \quad (4.2)$$

Die Intensität des auslösenden Lichts hat dagegen keinen Einfluss auf die Energie der austretenden Elektronen; sie bestimmt nur ihre Anzahl.

Die beiden für die Funkenkammer verwendeten Photomultiplier des Typs RCA 8575 detektieren Licht der Wellenlängen 260 nm bis 660 nm, wobei das Maximum des Ansprechvermögens bei ca. 358 nm liegt [Ori02].

4.2.1 Der Sekundärelektronenvervielfacher (SEV)

Der sich an die Fotokathode anschließende SEV hat die Aufgabe, die freien Elektronen zu sammeln und den entstehenden elektrischen Strom so zu verstärken, dass er nachgewiesen und verarbeitet werden kann.

Der SEV ist aus einer Anordnung von so genannten *Dynoden*, Elektroden aus einem Material mit hohem Sekundäremissionskoeffizienten (hier BeO), mit unterschiedlichem Potenzial aufgebaut.

Die aus der Fotokathode gelösten Elektronen werden über das elektrische Feld fokussierender Elektroden gebündelt und zur ersten Dynode beschleunigt. Beim Auftreffen werden dort weitere Elektronen aus ihren Bindungen gelöst. Das Sekundäremissionsvermögen, definiert als Anzahl der ausgelösten Elektronen je Anzahl der auftreffenden Elektronen, hängt stark von der kinetischen Energie der auftreffenden Elektronen ab. Schnelle Elektronen können mehrere Sekundärelektronen freisetzen, und man erhält ein verstärktes Signal (siehe Abb. 4.4). Dieses Prinzip wird beim SEV ausgenutzt, um einzelne Elektronen, d. h. geringe Lichtintensitäten bis herab zu einzelnen Photonen, nachzuweisen.

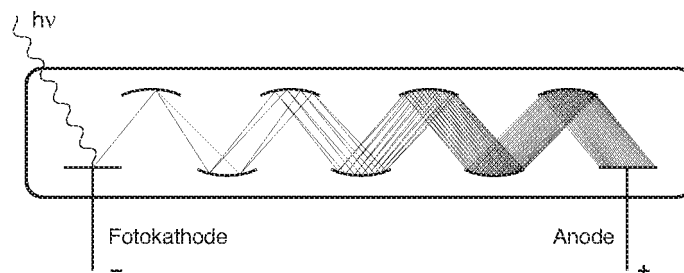


Abbildung 4.4: Schematische Darstellung der Stromverstärkung durch Auslösen von Sekundärelektronen.

Mit einer Anordnung von zwölf aufeinanderfolgenden Dynoden mit jeweils 150 – 200 V Potenzialdifferenz zwischen den Dynoden wird eine Vervielfachung der Elektronenzahl von 10^7 erreicht [Ori02]. Am Ende der Anordnung ist ein makroskopischer Strom von Elektronen entstanden. Die Gesamtladung an der Anode, die über einen Widerstand abfließt, kann als Spannung am Widerstand abgelesen werden; dieses ist das Ausgangssignal (siehe Kap. 4.2.2).

Die Geometrie der Dynoden muss so angelegt sein, dass ein hoher Prozentsatz der produzierten Elektronen an den jeweils folgenden Dynoden bzw. an der Anode gesammelt wird und Raumladungseffekte zwischen den Dynoden, vor allem aber zwischen letzter Dynode und Anode, vermieden werden (siehe Abb. 4.3).

4.2.2 Beschaltung der Base

Zu einem SEV gehört eine *Base*, in der sich ein Spannungsteiler befindet, der die anliegende Versorgungsspannung auf die einzelnen Dynoden verteilt. Dieser, hauptsächlich aus Widerständen aufgebaute, Teiler bestimmt also den Potenzialgradienten in der Röhre und damit auch das Zeit- und Verstärkungsverhalten.

Je nach Verwendungszweck werden passive oder aktive Basen benutzt. Eine passive Base ist im Wesentlichen aus einer Widerstandskette aufgebaut, in einer aktiven Base sind hingegen aktive Bauelemente, wie zum Beispiel Transistoren, integriert. Diese werden bei Experimenten mit hohen Zählraten benutzt, da hierbei die Spannung zwischen den einzelnen Dynoden nur noch durch aktive Bauelemente aufrecht erhalten werden kann.

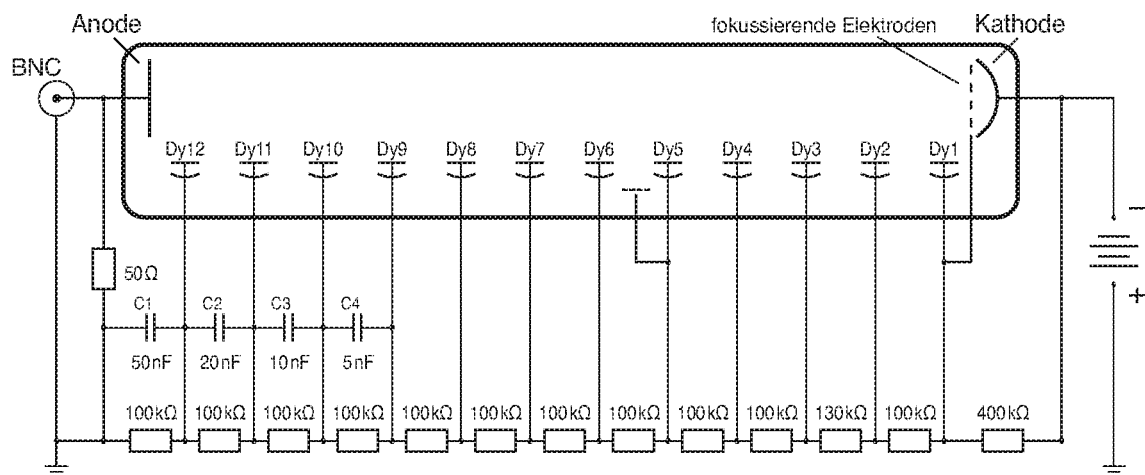


Abbildung 4.5: Ein typischer Spannungsteiler für den Photomultiplier des Typs RCA 8575.

Abb. 4.5 zeigt den Schaltplan eines typischen Spannungsteilers für den Photomultiplier des Typs RCA 8575. Der *Querstrom* durch die Widerstandskette muss groß sein gegen den Dynodenstrom, um Potenzialänderungen an den Dynoden während der Sekundärelektronenvervielfachung zu verhindern. Allgemein werden linearverstärkte Photomultiplier mit einer symmetrischen Spannungsteilung der Dynodenstufen, also gleichen Potenzialdifferenzen zwischen den Dynoden, betrieben. Das garantiert, dass der Emissionsfaktor jeder Dynode ungefähr gleich groß ist, d. h., dass pro Elektron im Mittel gleich viele Sekundärelektronen ausgelöst werden. Oft wird je nach Geometrie der Widerstand zwischen den ersten Dynoden etwas größer gewählt. So erhalten auch solche Elektronen eine genügend hohe kinetische Energie, die durch Photonen ausgelöst werden, die statt der Fotokathode direkt eine der ersten Dynoden treffen.

Die Kapazitäten C_1 bis C_4 haben die Aufgabe, bei einem großen Strom in den Vervielfachungskanälen des Multipliers die Potenziale der letzten Dynodenstufen zu stabilisieren (das Potenzial einer Dynode sinkt durch Ströme zwischen den Dynoden-zweigen). Da jede Dynode mit Widerstand Teil eines RC-Gliedes mit einer Resonanzfrequenz ist, sind die Kapazitäten außerdem zur Anode hin ansteigend, wodurch Resonanzen zwischen den Dynoden verhindert werden, die sich bis zur Anode fortpflanzen und somit der Pulsform als Schwingung überlagern könnten.

4.2.3 Messung der Stabilität der Pulsform zum Vergleich der Photomultiplier

Um die Funktionsweise der Photomultiplier (PM) zu überprüfen, ist es sinnvoll, eine Pulsformmessung vorzunehmen. Da es sich bei den hier verwendeten Photomultipliern um ältere Modelle handelt, von denen drei zur Verfügung standen, können auf diese Weise die zwei geeignetsten Photomultiplier-Base-Kombinationen gefunden werden. Mit dem Versuchsaufbau in Abb. 4.6 werden die Pulsformen der PM-Signale untersucht.

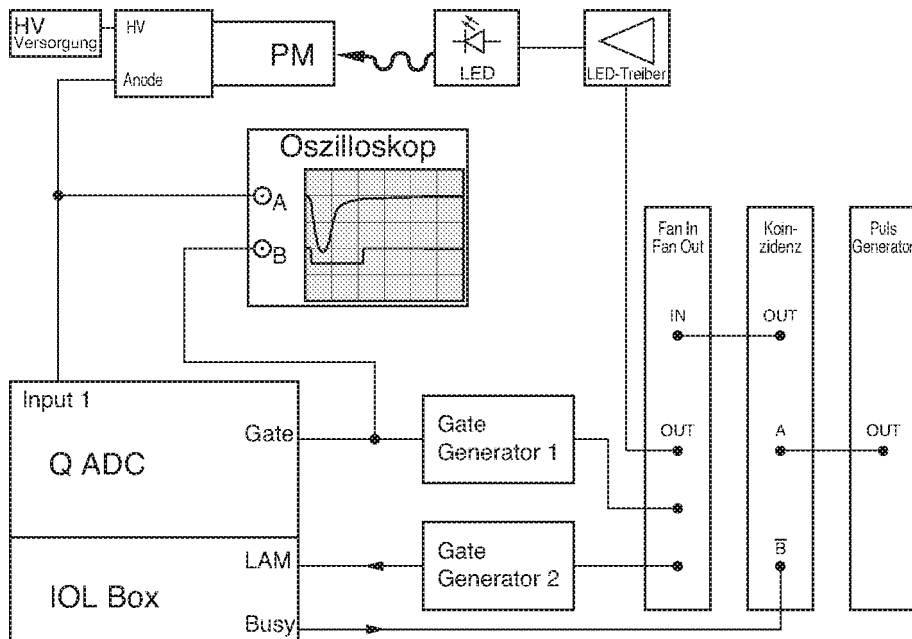


Abbildung 4.6: Aufbau zur Pulsformmessung der Photomultiplier-Signale.

Im Zentrum des Versuchsaufbaus steht eine Leuchtdiode (LED), die näherungsweise monochromatisches Licht der Wellenlänge $\lambda = 592 \text{ nm}$ emittiert, das von der zu untersuchenden Photomultiplier-Base-Kombination in ein elektrisches Signal um-

gesetzt wird. Zur Analyse geht das Anodensignal aus der Base an einen mit einem Rechner verbundenen *Q-ADC*, einen ladungssensitiven Analog-Digital-Converter.

Mit Hilfe eines LED-Treibers wird die LED gepulst. Die Frequenz, mit der dies geschieht, wird durch einen Pulsgenerator vorgegeben.

Um eine Signalaufnahme nur dann zu ermöglichen, wenn ein Ereignis stattgefunden hat, wird das Signal, das die LED ansteuert, gleichzeitig auf zwei *Gate Generatoren* gegeben. Der Erste hat die Aufgabe, ein Signal zu generieren, das auf den LAM (Look At Me)-Eingang einer Input/Output-Logik (IOL-Box) gegeben wird und mit einer eingestellten Verzögerung von etwa $6\ \mu\text{s}$ als Auslesesignal für die Datenerfassung dient.

Damit während der Datenerfassung keine weiteren Signale den Q-ADC erreichen, wird der Pulsgenerator von der IOL-Box für die Dauer der Q-ADC-Wandlung und -Auslese gesperrt, indem sie ein Signal B an eine Koinzidenzeinheit gibt, während der Rechner *busy* (beschäftigt) ist. Diese Koinzidenz-Einheit generiert das logische UND der Signale A des Pulsgenerators und $\bar{B}$ (nicht B) der IOL-Box; sie gibt also nur dann ein Signal aus, wenn der Rechner nicht busy ist.

Der zweite Gate Generator gibt gleichzeitig ein Gate (Zeitfenster) an den Q-ADC, der die Ausgangssignale der Photomultiplier-Base erfasst, so dass der PM nur ausgelesen wird, wenn ein LED-Puls stattfand.

Wenn ein Signal den Q-ADC erreicht und gleichzeitig das zweite Gate offen ist, wird die Ladung integriert und in einen Digitalwert umgewandelt, der ein Maß für die Ladungsmenge des Photomultiplier-Signals, also die Fläche unter dem Puls, ist. Ein im Rechner integrierter Vielkanalanalysator zählt die elektronischen Impulse entsprechend ihrer Höhe und füllt sie in ein Spektrum.

Bei der Auswahl der Photomultiplier anhand der aufgenommenen Spektren werden solche bevorzugt, die über einen langen Zeitraum stabil laufen. Da sie während des Versuchs immer durch die LED mit Licht gleicher Wellenlänge und Intensität bestrahlt wurden, sollte sich die Pulsform der Photomultiplier-Signale möglichst nicht ändern. Daher wurden aus den Spektren aller Photomultiplier-Base-Kombinationen die zwei herausgesucht, die die geringste mittlere Abweichung vom Maximum aufwiesen. Zur Verfügung standen jeweils drei Photomultiplier und drei Basen gleicher Bauart. Für jede der neun Kombinationen wurde ein Spektrum aufgenommen (siehe Abb. 4.7).

Mit diesem Versuchsaufbau konnte darüber hinaus die optimale Betriebsspannung für die Photomultiplier ermittelt werden, indem zunächst Spektren in Abhängigkeit von der Betriebsspannung aufgenommen wurden und die Spannung so eingestellt wurde, dass das Maximum jeweils im gleichen Kanal zu finden war. Das hat

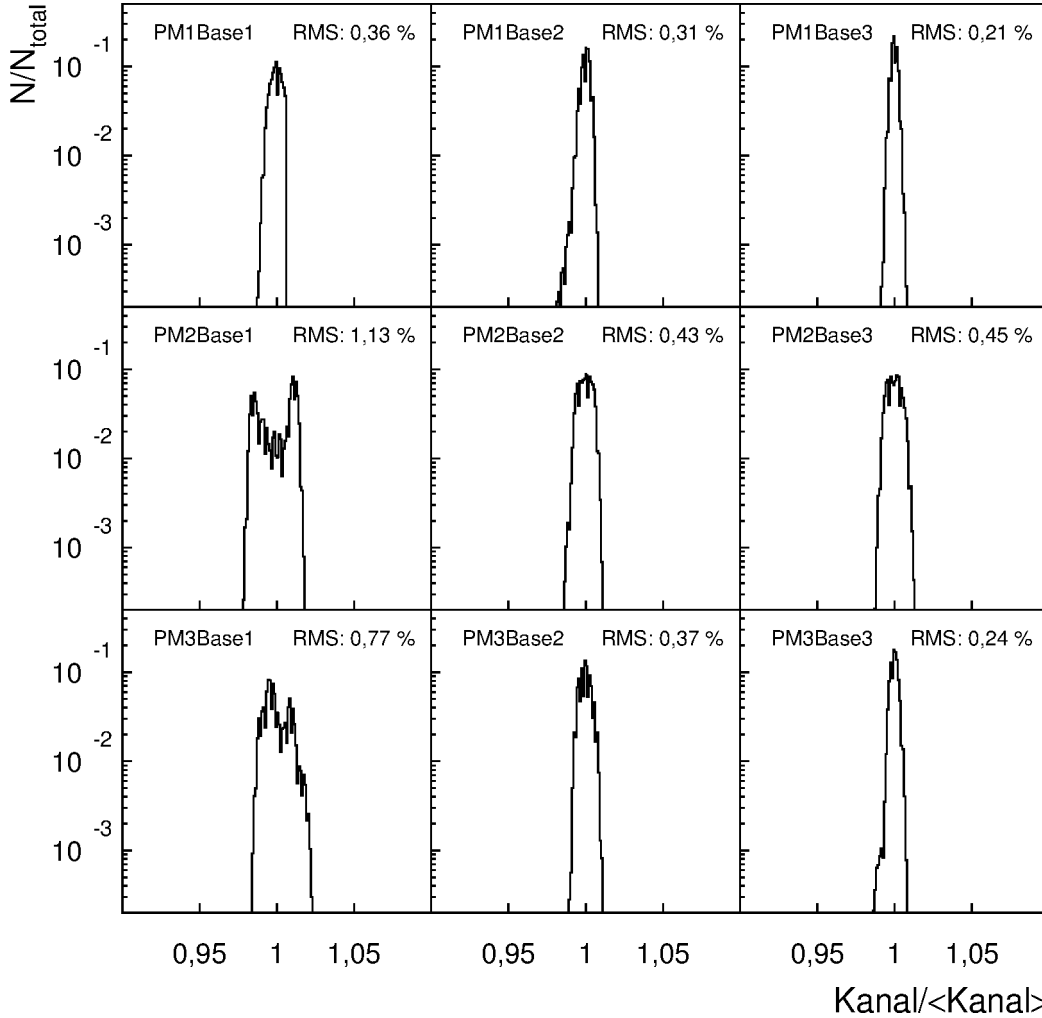


Abbildung 4.7: Pulsformspektren der verschiedenen Photomultiplier-Base-Kombinationen.

den Vorteil, dass die Photomultiplier mit dieser Spannung eine ähnliche Energieeichung besitzen. Allerdings war das lediglich mit einer Genauigkeit von ~ 50 Kanälen möglich.

Zur besseren Vergleichbarkeit sind die Spektren normiert, d. h. die Kanalnummern wurden jeweils durch den Mittelwert geteilt, so dass der Mittelwert jeder Distribution nun bei eins zu finden ist. Die Anzahl der Einträge der Spektren liegt im Bereich von zwei bis vier Millionen. Da auch dieser Wert nicht für jede Messung

identisch ist, wurden die Einträge ebenfalls normiert, also jeweils durch die Anzahl aller Einträge geteilt. Nun beträgt das Integral jedes Peaks eins und die Pulsformen können verglichen werden (siehe Abb. 4.7).

Die Spektren für die Kombinationen der Photomultiplier (PM1-PM3) mit den Basen (Base1-3) sind so angeordnet, dass in einer Zeile jeweils die Messungen mit gleichem Photomultiplier und in einer Spalte jeweils die mit gleicher Base zu finden sind. Der RMS (Root Mean Squared)-Wert ist die mittlere quadratische Abweichung vom Mittelwert (hier also von eins) und ist prozentual angegeben.

Nach einem Vergleich dieser Spektren wurden nach den Kriterien eines niedrigen RMS-Wertes und einer gleichmäßigen Pulsform die Kombinationen PM1/Base3 und PM3/Base2 ausgewählt, die später für die Triggerdetektoren der Funkenkammer dienen sollten. Zu beachten ist, dass die zwei besten Kombinationen unter der Voraussetzung gefunden werden mussten, dass jedes Bauteil nur einmal verwendet werden kann.

4.3 Die Triggerelektronik

Die Signale der Photomultiplier müssen nun so weiterverarbeitet werden, dass nur bei gleichzeitigem Ansprechen beider Paddles ein Triggersignal an die Zündansteuerung weitergegeben wird.

Die Definition dieses Triggers ist also das logische UND aus den Signalen der Photomultiplier. Abbildung 4.8 zeigt den Aufbau des Triggers inklusive der elektronischen Auslekette der beiden Triggerdetektoren. Zur Erzeugung des Triggersignals kommen eine Koinzidenzeinheit und Diskriminatoren zum Einsatz.

Der *Constant-Fraction-Discriminator* (CFD) wandelt ein Eingangssignal in ein digitales Normsignal um, sobald dessen Amplitude einen vorher eingestellten Anteil an der Gesamtamplitude überschreitet: Das Eingangssignal wird um eine feste Zeit verzögert und außerdem gedämpft und invertiert. Die beiden so erhaltenen Signale werden voneinander subtrahiert, und die Diskriminatorschwelle resultiert aus einem Nulldurchgang dieses Differenzsignals, dessen Zeit weitgehend unabhängig von der Amplitude des ursprünglichen Signals ist. Ein CFD-Modul besteht aus mehreren diskreten CFDs, die analogen Signale der beiden Photomultiplier werden jeweils auf einen CFD gegeben.

Die logische UND-Verknüpfung dieser Signale erfolgt in einer Koinzidenzeinheit. Sie gibt ein Signal aus, wenn an beiden Eingängen *gleichzeitig* ein CFD-Signal anliegt. Von gleichzeitig darf gesprochen werden, da sich die Myonen nahezu mit Lichtgeschwindigkeit bewegen und die zeitliche Differenz der Signale klein ist gegen die Breite der Normsignale.

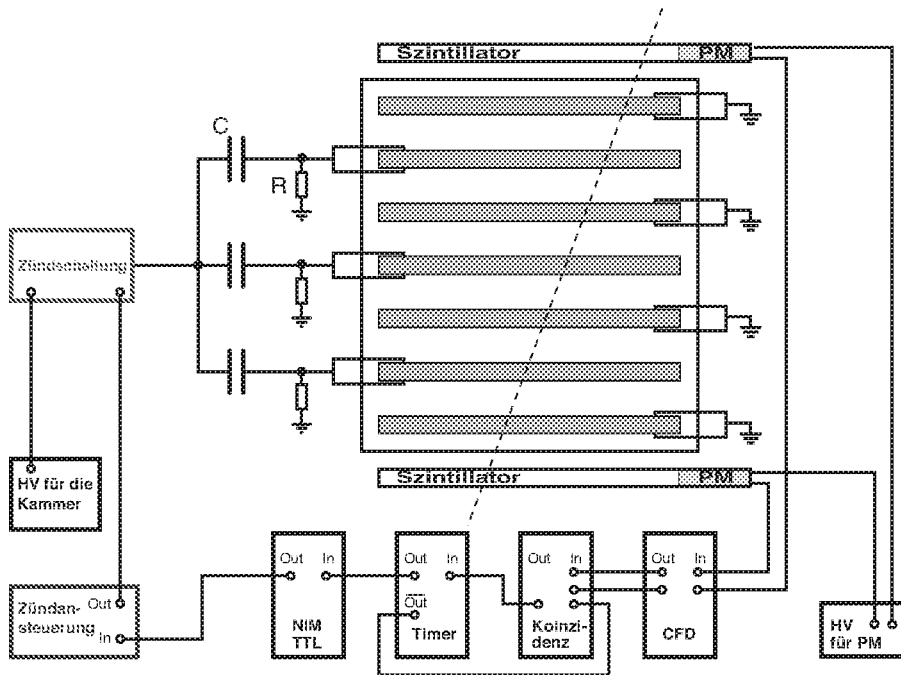


Abbildung 4.8: Der Aufbau des Triggers inklusive der Ausleseelektronik.

Das Ausgangssignal der Koinzidenzeinheit könnte nun direkt für die Zündansteuerung genutzt werden, allerdings muss gewartet werden, bis der Detektor nach einem vorherigen Teilchendurchgang wieder sensitiv ist. Diese Totzeit ergibt sich zum größten Teil aus der Aufladezeit der Kondensatoren. An den Platten der Kammer befinden sich 18 Kondensatoren der Kapazität $C = 1 \text{ nF}$, die nach einer Funkenentladung über einen $1,7 \text{ M}\Omega$ -Widerstand wieder aufgeladen werden (siehe Abb. 5.1). In der Zündansteuerung befindet sich ein Kondensator mit $C = 100 \text{ nF}$, dessen Aufladewiderstand sich aus zwei parallel geschalteten Widerständen mit $R_1 = 330 \text{ k}\Omega$ und $R_2 = 67 \text{ k}\Omega$ zusammensetzt (siehe Abb. 5.3). Innerhalb der größeren der beiden Aufladezeiten sollten alle Kondensatoren wieder aufgeladen sein.

Die Aufladefunktion eines Kondensators ist eine beschränkte exponentielle Wachstumsfunktion $U(t) = U_{\text{max}} \cdot (1 - e^{-\frac{t}{\tau}})$ mit der angelegten Spannung U und der Zeitkonstanten τ . Die Ladezeit ist nur von der Kapazität C des Kondensators und dem Widerstand R abhängig. Daher wird das Produkt aus C und R als Zeitkonstante τ festgelegt. Innerhalb einer Zeitkonstanten lädt sich ein Kondensator auf $1 - e^{-1} = 63 \%$ der angelegten Spannung auf. Nach $0,69\tau$ hat er ca. 50% seiner endgültigen Spannung erreicht und nach fünf Zeitkonstanten ist ein Kondensator fast vollständig aufgeladen (zu mehr als 99%). Zunächst sind die einzelnen Zeitkonstanten zu betrachten:

Zeitkonstante der 18 Kondensatoren an den Platten der Kammer:

$$\tau_1 = 18 \cdot 1,7 \text{ M}\Omega \cdot 1 \text{ nF} = 30,6 \text{ ms}, \quad (4.3)$$

Zeitkonstante des Kondensators in der Zündansteuerung:

$$\tau_2 = 55,7 \text{ k}\Omega \cdot 100 \text{ nF} = 5,6 \text{ ms}. \quad (4.4)$$

Daraus folgt, dass nach $5\tau_1 = 153 \text{ ms}$ alle Kondensatoren aufgeladen sind. Dieser Wert wird als erste Abschätzung der Totzeit der Funkenkammer genommen.

Es ist nun die Aufgabe eines Timers, das erneute Anlegen der Hochspannung an die Platten während der Totzeit zu verhindern. Bekommt der Timer ein Signal aus der Koinzidenzeinheit, läuft eine zuvor eingestellte Uhr ab. Währenddessen wird ein invertiertes OUT-Signal auf die Koinzidenzeinheit gegeben und weitere Signale werden blockiert. Alternativ kann auch ein OUT-Signal des Timers in Antikoinzidenz zu den beiden Signalen der CFDs geschaltet werden. Darüber hinaus gibt der Timer das ankommende Signal weiter, welches durch einen NIM-TTL-Wandler in ein positives TTL-Signal mit einer Amplitude von 5 V gewandelt wird (siehe kurze Erläuterung im Anschluss). Dieses ist das *Triggersignal*, welches die Schaltlogik für die Hochspannung ansteuert.

Mit dem Timer ist es möglich, die gesamte Totzeit der Funkenkammer abzuschätzen: Die Wartezeit kann bis auf 0,45 s herunter gesetzt werden und es sind immer noch nach jedem Zündvorgang Funkenentladungen zu beobachten. Erst unter dieser Zeit nimmt der Anteil der Funkenentladungen an den getriggerten Ereignissen ab.

Für den Dauerbetrieb wird eine Wartezeit von 2 s gewählt, die also wesentlich größer als die Totzeit ist. Eine höhere Rate wäre nicht wünschenswert, da durch Funkenerosion eine starke Abnutzung der Zündkerze und der gegenüberliegenden Elektrode in der Zündschaltung eintritt (siehe Kap. 5.1).

Eine kurze Erläuterung von NIM- und TTL-Logik

Es gibt verschiedene Arten, logische Schaltungen zu realisieren, die so genannten „Logikfamilien“. Die Bezeichnung TTL bedeutet *Transistor-Transistor-Logik*. Die TTL-Familie wurde 1964 eingeführt und entwickelte sich bald zur verbreitetsten Logikfamilie. Sie war im Wesentlichen eine Weiterentwicklung der bis dahin eingesetzten DTL- (Dioden Transistor Logik) und RTL- (Widerstands (engl.: Resistor) Transistor Logik) Familien, deren Schaltungszustand von einem *aktiven Bauelement* (Transistor) zu einem *passiven Bauelement* (Diode) bestimmt wurde. Die Schaltzeiten, die damit erreicht wurden, sind entsprechend langsam. Bei der TTL-Familie werden beide Schaltzustände durch aktive Bauelemente (deshalb Transistor-Transistor-Logik)

realisiert. Dies ermöglicht einen schnelleren Wechsel der logischen Zustände 0 und 1, so dass Schaltzeiten von 10 ns erreicht werden. Die aktiven Elemente (Transistoren) werden in die *Sättigung* geschaltet. So wird der Low-Pegel des Ausgangs durch die Sättigungsspannung von Kollektor zu Emitter bestimmt und liegt bei einer Spannung von 0,2 V (siehe dazu [HBG94]). Der High-Pegel liegt bei den hier verwendeten TTL-Bauteilen bei +5 V.

Die Abkürzung NIM steht für *Nuclear Instrumentation Modules* und ist ein Standard, der ausschließlich in der Kern- und Teilchenphysik Anwendung findet. NIM-Module beinhalten analoge und digitale Instrumente. Analoge Signale tragen zusätzlich Informationen in ihrer Amplitude und Form, so dass sie kontinuierlich in Form und Höhe variieren. Digitale Signale können wieder zwei Zustände (logisch 0 und logisch 1) annehmen. Es existieren zwei Standards: die so genannte *slow-positive Logic* und die *fast-negative Logic*. Die Signale der ersten Logik sind Signale mit relativ langsamer Anstiegszeit in der Größenordnung von hunderten Nanosekunden und mehr. Sie haben positive Polarität und finden ihre Anwendung in langsamen Detektorsystemen. Die zweite Logik (oft als *NIM-Logik* bezeichnet) ist Strom basierend und verwendet extrem schnelle Signale mit Anstiegszeiten von etwa einer Nanosekunde und vergleichbaren Signaltbreiten. Solche Signale werden vor allem in der Hochenergiephysik benötigt, wo hohe Zählraten oder schnelles Timing erwünscht sind. Die Leitungen müssen mit einem Widerstand von 50 Ω abgeschlossen sein (siehe Kap. 5.2.4), die entsprechenden Spannungslevel liegen bei 0 V (logisch 0) und -0,8 V (logisch 1). Im Gegensatz zu slow-positive Signalen können fast-NIM Signale auf Grund der Stromlogik über relativ lange Kabel verlustfrei übertragen werden (siehe dazu [Leo87]).

4.4 Die geometrische Akzeptanz und Effizienz des Triggers

Die Rate der Triggersignale und somit die maximale Rate, mit der Funkenentladungen in der Funkenkammer stattfinden, ergibt sich zum einen aus der *geometrischen Akzeptanz* der Triggerdetektoren und darüber hinaus aus ihrer *Effizienz*, d.h. der Wahrscheinlichkeit, auf einen Teilchendurchgang zu reagieren.

Desweiteren wird die Rate der Funkenentladungen durch die vom Timer eingestellte Totzeit begrenzt (siehe Kap. 4.3).

4.4.1 Die geometrische Akzeptanz der Triggerdetektoren

Die Gesamtrate der Koinzidenzen wird zunächst von der geometrischen Akzeptanz der Triggerdetektoren beeinflusst, da sie die Rate der Teilchen bestimmt, die beide Szintillatoren passieren können. Die Szintillatoren haben eine Fläche von jeweils $72 \times 20 \text{ cm}^2$ und sind in einem Abstand von 50 cm mit einem Versatz von 5 cm angebracht. Wird die aus der Theorie zu erwartende Myonenrate mit einer Verteilung proportional zu $\cos^2 \theta$ berücksichtigt (siehe Kap. 2.3), so lässt sich die erwartete Koinzidenzrate bei einer Nachweiswahrscheinlichkeit von 100 % ermitteln. Die Abschätzung der Rate bei gegebenen Parametern erfolgt mit einer Monte Carlo-Simulation. Daraus ergibt sich eine Koinzidenzrate von 9,95 Koinzidenzen pro Sekunde.

4.4.2 Die Effizienz der Triggerdetektoren

Die gemessene Rate wahrer Koinzidenzen beträgt $3,58 \pm 0,6$ Koinzidenzen pro Sekunde (siehe Kap. 4.4.3). Nach der geometrischen Abschätzung durch die Monte Carlo-Simulation löst also etwa jeder dritte Teilchendurchgang eine Koinzidenz aus, was einer Effizienz der Triggerdetektoren von jeweils etwa 60 % entspräche. Allerdings wird die Anzahl der gemessenen Koinzidenzen durch Diskriminator-Schwellen herabgesetzt, die nicht nur ein elektronisches Rauschen unterdrücken sondern auch Signale von Myonen, die nicht genügend Energie im Szintillator deponieren. Die Monte Carlo-Simulation berücksichtigt weder diese Schwellen noch andere Effekte wie die der Lichtkopplung in den Paddles oder der Elektronik.

Optimierung der Effizienz der Triggerdetektoren

Die Effizienz der Triggerdetektoren, d. h. die Wahrscheinlichkeit, dass sie auf einen Teilchendurchgang reagieren, setzt sich aus der Effizienz der Szintillatoren und der Lichtkopplung in den Paddles bzw. zwischen Paddle und Photomultiplier zusammen.

Es ist sehr schwierig, die Effizienz exakt zu bestimmen, da sie als Verhältnis zwischen Eingangs- und Ausgangsverteilung definiert ist, die Eingangsverteilung allerdings nicht exakt bestimmbar ist, weil keine definierte Myonenquelle zur Verfügung steht.

Die folgenden Messungen wurden an den Triggerdetektoren vorgenommen, als diese bereits fest an der Kammer montiert waren. Die Anodensignale der Photomultiplier werden mit Hilfe eines in einem Rechner integrierten Vielkanalanalysators nach der Amplitudenhöhe sortiert und in ein Spektrum gefüllt.

Zunächst soll mit diesem Versuchsaufbau überprüft werden, ob von außen Licht in die Paddles gelangt, welches unerwünschte Signale im Photomultiplier erzeugt,

Oberer Triggerdetektor	Messdauer	Zählrate
dunkler Raum	1 min	13 379
mit Beleuchtung	1 min	3 238 745
nach Ausbesserungen mit Beleuchtung	1 min	13 075
Unterer Triggerdetektor	Messdauer	Zählrate
dunkler Raum	1 min	56 144
mit Beleuchtung	1 min	3 982 145
nach Ausbesserungen mit Beleuchtung	1 min	58 963

Tabelle 4.1: Messungen zur Lichtundurchlässigkeit der Triggerdetektoren.

die zu mehr zufälligen Koinzidenzen führen. Dazu wird zunächst die Signalrate des oberen Detektors im dunklen Raum und zum Vergleich bei eingeschalteter Deckenbeleuchtung bestimmt. Die Ergebnisse dieser Messungen sind in Tab. 4.1 dargestellt.

Zunächst fiel auf, dass die Zählrate im dunklen Raum erheblich kleiner war als die Zählrate bei eingeschalteter Deckenbeleuchtung. Daher gab es vermutlich noch Stellen, an denen Photonen in den Szintillator oder Lichtleiter eindringen konnten. Um diese Stellen zu finden, wurde die Messung im dunklen Raum wiederholt, wobei die Paddles mit einer kleinen Taschenlampe abgefahren wurden. Gleichzeitiges Betrachten der sich füllenden Spektren auf dem Monitor des Rechners lieferte das Ergebnis, dass die Detektoren Schwachstellen bezüglich der Lichtundurchlässigkeit an den Kopplungsstellen zwischen Lichtleiter und Photomultiplier aufwiesen. Nach erneutem Abkleben dieser kritischen Stellen sank die Rate auch bei eingeschalteter Beleuchtung bis auf einen Wert, der mit dem im dunklen Raum vergleichbar war. Ein sehr ähnliches Ergebnis lieferte die Messung mit dem unteren Detektor (siehe Tab. 4.1).

Zusammenfassend lieferten diese Messungen die Erkenntnis, dass die Lichtundurchlässigkeit der Triggerdetektoren sehr empfindlich ist und regelmäßig überprüft werden sollte.

Mit dem gleichen Versuchsaufbau lässt sich die Lage des Maximums der Signalarhöhe der Photomultiplier in Abhängigkeit von der Position einer radioaktiven Quelle (^{137}Cs) auf dem Szintillator bestimmen. Die Quelle emittiert Elektronen und Photonen, welche im Szintillator Signale auslösen können. Diese Messungen sollten die Abhängigkeit der Pulshöhe vom Ort des Teilchendurchgangs und somit eine eventuell

nicht gleichmäßige Effizienz der Triggerdetektoren untersuchen. Das erste Ergebnis dieser Messungen war, dass der obere Photomultiplier erst nach einer Zeit von ca. 20 Minuten nach dem Anlegen der Hochspannung sinnvolle Signale liefert. Dieses Verhalten fiel bei der ersten Messung der Pulsform (vergleiche Kap. 4.2.3) nicht auf, da es sich dort ausschließlich um Langzeitmessungen handelte. Darüber hinaus misst dieser Photomultiplier bei höherer anliegender Versorgungsspannung eine niedrigere Rate als der andere, was auf Probleme mit der Base hindeutet. Ein späterer Austausch dieses Photomultipliers wäre sinnvoll, ist aber bislang nicht vorgenommen worden, da die Effizienz für die Zwecke der Funkenkammer, die ausschließlich als Demonstrationsobjekt dient und jeweils für längere Zeit in Betrieb genommen wird, ausreicht.

Die Spektren zeigten die Häufigkeiten von Signalen bestimmter Höhe und es gab ein sehr kleines lokales Maximum, das mit sinkendem Abstand zwischen Quelle und dem Ende des Szintillators an der Lichtleiterseite zu größeren Pulshöhen wanderte. Dieses lokale Maximum war allerdings so schwach ausgebildet, dass es für mittlere Positionen der Quelle ganz in einem wesentlich größeren Peak (wahrscheinlich elektronischer Herkunft) verschwand. Aus diesem Grund ließ sich mit diesem Versuch keine ausgeprägte Abhängigkeit der Effizienzhöhe vom Ort der Detektion feststellen.

4.4.3 Die Koinzidenzrate

Ohne Timer ist die maximal mögliche Anzahl der Funkenentladungen gleich der Koinzidenzrate der Triggerdetektoren. Die gemessene Koinzidenzrate ist von mehreren Variablen abhängig. Eine Vergrößerung der Hochspannung an den Photomultipliern vergrößert auch ihre Zählrate, weil dadurch die Verstärkung des Photomultipliers erhöht wird und auch Teilchen geringerer Energien ein Ausgangssignal liefern. Insgesamt überschreiten dadurch mehr Ereignisse die Diskriminatorschwellen, durch Erhöhung der Schwellen wird die Koinzidenzrate wiederum gemindert.

Die Gesamtrate der Koinzidenzen N_{ges} ist gleich der Summe aus der Rate *wahrer Koinzidenzen* $N_{1,2}$ und der Rate *zufälliger Koinzidenzen* N_{zuf} . Eine wahre Koinzidenz wird gemessen, wenn beide Szintillatoren dasselbe Teilchen detektieren. Darüber hinaus werden zufällige Koinzidenzen gemessen, die daraus resultieren, dass die Szintillatoren auch voneinander unabhängige Ereignisse gleichzeitig detektieren. Diese können zum Beispiel zwei unterschiedliche Myonen sein, aber auch Elektronen oder Photonen. Das Auslösen eines Triggersignals durch zufällige Koinzidenzen ist für die Funkenkammer unerwünscht, da sie keine Funkenstrecken durch die ganze Kammer zur Folge haben. Die Diskriminatoren verhindern die Verarbeitung eines Großteils des Untergrundes in den Photomultiplier-Signalen, so dass fast ausschließ-

lich die Signale hochenergetischer Myonen die Koinzidenzeinheit erreichen. Dadurch wird die Rate der zufälligen Koinzidenzen stark reduziert.

Zunächst sind die Einzelraten der Photomultiplier bei den eingestellten Bedingungen zu betrachten, die in Tab. 4.2 zu finden sind.

	PM oben	PM unten
HV (V)	2317	2087
Diskriminatorschwelle (mV)	107,6	111,3
Zählrate mit Diskriminator (Hz)	$21,8 \pm 0,6$	$29,9 \pm 0,7$

Tabelle 4.2: Die Einzelzählraten der Photomultiplier.

Der obere Photomultiplier hat eine durchschnittliche Zählrate N_1 , die Signale haben nach der Diskriminierung eine Breite ρ_1 . Analog erreicht der untere Photomultiplier eine durchschnittliche Zählrate N_2 mit einer Signalbreite nach Durchlaufen des Diskriminators von ρ_2 . Die Einzelraten der Photomultiplier N_1 und N_2 sind generell wesentlich größer als die Gesamtrate der Koinzidenzen und somit natürlich auch größer als N_{zuf} . Genähert gilt für die Rate der zufälligen Koinzidenzen [Eva55]:

$$N_{\text{zuf}} \approx N_1 N_2 (\rho_1 + \rho_2). \quad (4.5)$$

Die beiden hier benutzten Diskriminatoren sind vom gleichen Typ und geben in Form und Amplitude identische NIM-Signale aus. Die Breite der Signale wurde mit einem Oszilloskop bestimmt, und es gilt: $\rho_1 \approx \rho_2 \approx 60$ ns. Daraus ergibt sich nach Gleichung (4.5) eine Rate zufälliger Koinzidenzen von etwa $6 \cdot 10^{-5}$ Hz.

Zur Überprüfung wurden die Gesamtraten der Koinzidenzen gemessen: Wird jetzt auf eines der Photomultiplier-Signale eine Verzögerung (Delay) gegeben, so sinkt die Rate wahrer Koinzidenzen, so dass bei einem ausreichend großen Delay nur noch zufällige Koinzidenzen gemessen werden. Die Größe dieses Delays ist von der Breite der Signale aus den Diskriminatoren abhängig. Diese sind rechteckförmige NIM-Signale mit einer Breite von etwa 60 ns, wobei die Flanken nahezu als senkrecht angesehen werden können.

Aus Tab. 4.3 folgt, dass die Gesamtrate der Koinzidenzen bis zu einer bestimmten Verzögerung eines der Photomultiplier-Signale konstant bleibt. Wird das Signal des oberen Photomultipliers um 60 ns und mehr verzögert, nimmt die Rate sehr rasch ab, und schon bei einer Verzögerung von 75 ns liegt die Rate bei unter einem Promille (dies ergibt sich aus der Messdauer von 20 Minuten). Wird stattdessen das Signal des unteren Photomultipliers verzögert, so sinkt die Rate schon bei einer Verzögerung von 60 ns auf einen so niedrigen Wert, dass innerhalb von 20 Minuten keine

Verzögerung (ns)	gesamte Koinzidenzrate (Hz)	
	Verz. am oberen PM	Verz. am unteren PM
0	$3,58 \pm 0,03$	$3,58 \pm 0,03$
9	$3,65 \pm 0,11$	$3,62 \pm 0,11$
14	$3,64 \pm 0,11$	$3,32 \pm 0,11$
24	$3,43 \pm 0,11$	$3,17 \pm 0,10$
36	$3,14 \pm 0,10$	$2,61 \pm 0,09$
44	$3,16 \pm 0,10$	-
46	-	$1,10 \pm 0,06$
51	-	$0,28 \pm 0,03$
56	$2,93 \pm 0,10$	$0,02 \pm 0,01$
71	$1,91 \pm 0,06$	-
72	$2,80 \pm 0,10$	0
74	$0,97 \pm 0,06$	-
79	$0,12 \pm 0,02$	-
89	0	-

Tabelle 4.3: Messung der Koinzidenzrate mit verschiedenen Verzögerungen.

Koinzidenzen mehr registriert wurden. Die Werte für die Verzögerungen ergeben sich aus einem durch ein NIM-Modul erzeugten Delay und einer weiteren Verzögerung durch unterschiedliche Kabellängen zur Beschaltung dieser Module. Aufgetragen ist hier die Summe der Verzögerungen.

Abbildung 4.9 zeigt die Koinzidenzrate in Abhängigkeit von der Verzögerung auf einem der Signale für zwei sich überlagernde ideale Rechtecksignale (grauer Bereich) und für die hier gemessenen Daten (Punkte).

Die Tatsache, dass die gemessenen Koinzidenzraten nicht so steil abfallen wie es im idealen Fall zu erwarten wäre, liegt an der nicht ganz idealen Rechteckform der Diskriminator-Signale. Sie überlagern sich über einen Verzögerungsraum von etwa $2 \cdot 60$ ns mit voller Amplitude, und der Überlapp sinkt dann mit der Flanke. Im Idealfall (logisches UND zweier Rechtecksignale) gäbe es einen senkrechten Abfall der Rate im Diagramm bei einer bestimmten Verzögerung, die der Breite des Signals entspricht.

Ein weiterer Effekt lässt sich Abb. 4.9 entnehmen: die Verteilung der Koinzidenzraten ist nicht symmetrisch um null wie die idealisierte Antwort zweier Rechtecksignale ohne zusätzliche Verzögerung. Das liegt daran, dass hier auch die Signale

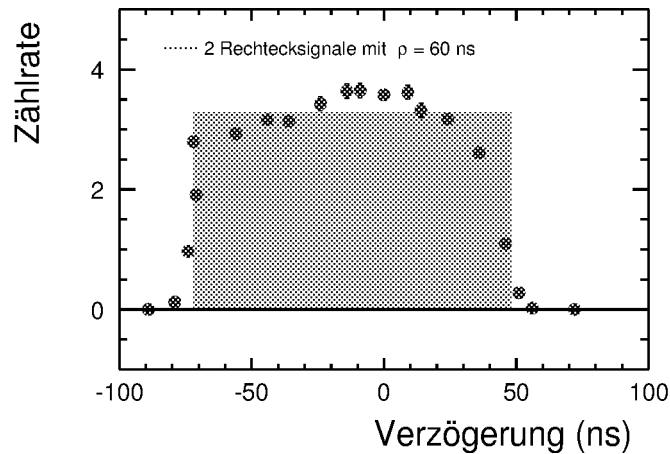


Abbildung 4.9: Die Koinzidenzrate als Funktion der Verzögerung eines der Signale für zwei sich überlagernde Rechtecksignale schematisch (ausgefüllte Fläche) und gemessen (Punkte). Eine negative Verzögerung bedeutet, dass das Signal des unteren Photomultipliers verzögert wurde, eine positive ist eine Verzögerung des oberen Photomultiplier-Signals.

einer wahren Koinzidenz ohne Verzögerung nicht genau aufeinander liegen, weil die Signale auf dem Weg vom Photomultiplier bis zum Diskriminator nicht die gleiche Verzögerung erfahren, da sie unterschiedliche Entfernungen zu überwinden haben. Das Kabel zum unteren Photomultiplier ist kürzer. Wird nun das Signal des oberen Photomultipliers verzögert, so liegen die Signale für kleine Verzögerungen sogar noch besser übereinander als ohne. Hingegen wird bei einer Verzögerung des unteren Photomultipliers die ohnehin schon vorhandene Laufzeitdifferenz vergrößert.

Die Differenz der Verzögerungen, bei denen die Koinzidenzrate in unserer Messzeit auf null sinkt, beträgt in beiden Richtungen etwa 10 ns. Dieser Wert für die Verschiebung der Signale zueinander stimmt bei einer Versorgung durch ein ca. 2 m längeres Kabel des unteren Photomultipliers gut mit dem angenommenen Wert einer Verzögerung von 5 ns/m für ein Kabel überein.

Je höher die Diskriminatorschwellen liegen, desto größer wird der Anteil wahrer Koinzidenzen. Das Ergebnis dieser Messungen ist, dass für diese Einstellungen der Diskriminatoren die Rate der zufälligen Koinzidenzen unter einem Promille liegt. Dieses Ergebnis wird durch die Beobachtung bestätigt, dass nach einer Koinzidenz, d. h. dem Triggern der Kammer, praktisch immer eine Funkenspur zu erkennen ist. Mit den hier gewählten Einstellungen wird eine Gesamtkoinzidenzrate und somit auch Rate wahrer Koinzidenzen von $(3,58 \pm 0,6)$ Hz gemessen.

Dieser Wert für die Koinzidenzrate wurde nach der ersten Inbetriebnahme der Funkenkammer nicht erreicht: Die Rate betrug zunächst lediglich 0,1 Hz. Erst nach einmaligem, aber extremem, Hochfahren der am oberen Photomultiplier anliegenden Spannung auf über 2,5 kV kam es zu einer sofortigen Erhöhung der Koinzidenzrate auf 3,6 Hz, die jetzt bei den vorherigen Einstellungen der Schwellen und der Hochspannung stabil gemessen wird. Dies ist ein weiteres Indiz für Probleme mit der Base des oberen Photomultipliers, die auf Dauer ersetzt werden sollte.

5. Das Zündsystem

Um Funkenentladungen am Ort der Primärionisation zu erzeugen, muss in kürzester Zeit eine Hochspannung von ca. 5 kV (siehe Kap. 3.1) an die Platten der Funkenkammer gelegt werden. Als Schalter dient eine Funkenstrecke, die über eine Hilfsfunkenstrecke gezündet wird. Die Zündschaltung umfasst den Aufbau für diese Funkenstrecken inklusive einer Zündspule (siehe Kap. 5.1). Die Zündansteuerung dient als schnelles Schaltsystem, das nach einem Triggersignal einen Strom an die Zündspule legt und somit die Hilfsfunkenstrecke auslöst (siehe Kap. 5.2).

5.1 Die Zündschaltung

Die Zündschaltung hat die Aufgabe, innerhalb einer sehr kurzen Zeit die Hochspannung an die Platten der Funkenkammer zu leiten. Die Zeitspanne zwischen Teilchendurchgang und Anlegen der Hochspannung sollte nicht größer als $1\ \mu\text{s}$ sein, damit sich die in der Kammer entstandenen Ladungsträger nicht zu weit von der Teilchenspur entfernen (siehe Kap. 3.2). Dazu werden insgesamt 18 Kondensatoren der Kapazität $1\ \text{nF}$ an jede zweite Platte angeschlossen und mit der Hochspannungsquelle verbunden (siehe Abb. 5.1).

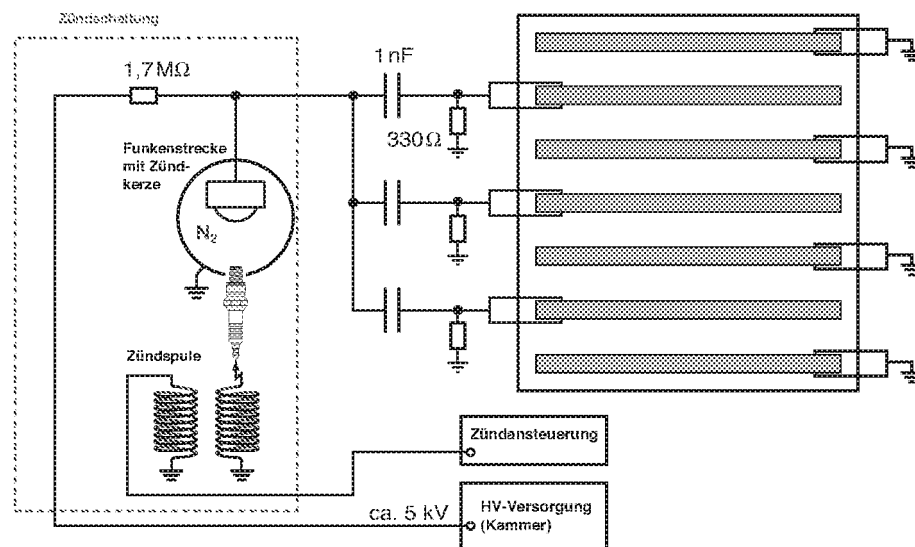


Abbildung 5.1: Die Zündschaltung.

Die Kondensatoren werden über einen Widerstand von $1,7\text{ M}\Omega$ aufgeladen. Nach dem Ladevorgang wirken die Kondensatoren als unendlich große Widerstände und es fließt kein Strom mehr.

Die Hauptfunkenstrecke (Hochspannungsschalter) zwischen der Masseelektrode einer Zündkerze und einer Messingelektrode, befindet sich zwecks schnellerer Zündung innerhalb eines mit Stickstoff gefüllten Gehäuses, wobei der Spitzenstrom einer Entladung einen Wert größer als 1000 A erreichen kann. Über die Zündspule wird eine Spannung an die Zündkerze angelegt, deren Kopf in das Gehäuse hineingedreht ist (siehe Kap. 5.1.1). Auf der gegenüberliegenden Seite befindet sich die Messingelektrode, die mit der Hochspannungsquelle verbunden ist. Die an der Mittelelektrode der Zündkerze entstandenen Ladungsträger driften in Richtung Elektrode, bis eine leitende Verbindung entsteht und sich ein Funken ausbildet. In diesem Moment sind die Kondensatoren geerdet. Durch die Potenzialänderung an den Kondensatoren wird ihre Ladung an die Platten der Funkenkammer abgegeben.

5.1.1 Die Zündkerze

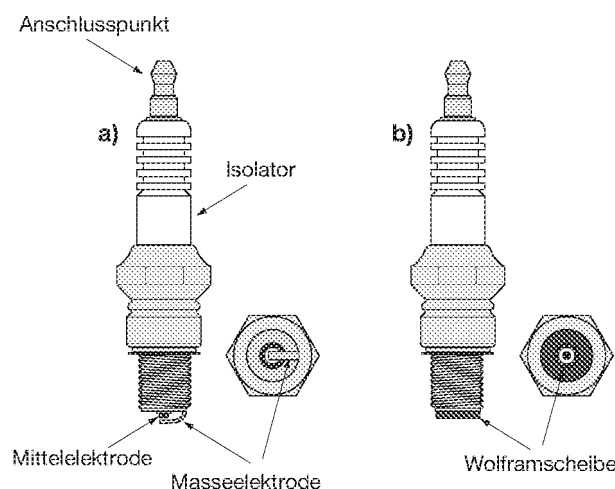


Abbildung 5.2: Eine handelsübliche Zündkerze a) und mit modifizierter Elektrode b).

Die Zündkerze hat die Aufgabe mit der in der Zündspule erzeugten Spannung, eine Hilfsfunkenstrecke zu zünden. Bestandteile der Zündkerze sind ein Metallmantel und ein Isolator aus Keramik (siehe Abb. 5.2a). Im Inneren des Isolators befindet sich die Mittelelektrode, an der das Zündkabel angeschlossen wird. Die Masseelektrode, ein gebogener Metallhaken an der Vorderseite, ist mit dem Gehäuse der Zündkerze verbunden. Die Mittelelektrode geht durch die gesamte Zündkerze hindurch und

steht am Ende der Masseelektrode gegenüber. Der Abstand zwischen diesen Elektroden muss abhängig von der angelegten Spannung gewählt werden.

Die speziell für die Funkenkammer modifizierte Zündkerze besitzt als Masseelektrode statt des Metallhakens eine Wolframscheibe (siehe Abb. 5.2b). Diese Wolframscheibe liegt als Ring um die Mittelelektrode und hat auf Grund ihrer vergrößerten Fläche den Vorteil, dass der Funke nicht mehr durch die Elektrode verdeckt wird, sondern in das Kammervolumen hinein reicht, was ein schnelleres Zünden der Hauptfunkenstrecke bewirkt. Außerdem ist die modifizierte Masseelektrode wesentlich resistenter gegen Funkenerosion als handelsübliche Zündkerzen, was den Wartungsaufwand der Funkenkammer minimiert.

5.2 Die Zündansteuerung

Die Zündkerze bezieht ihren Zündimpuls aus der Zündansteuerung, welche nach Eingang eines Triggersignals den notwendigen Spannungsimpuls liefert. Die Zündansteuerung ist der Schalter, der bei Eingang eines Triggersignals den Strom durch die Zündspule auslöst, um die Hilfsfunkenstrecke zu zünden (siehe Abb. 5.3).

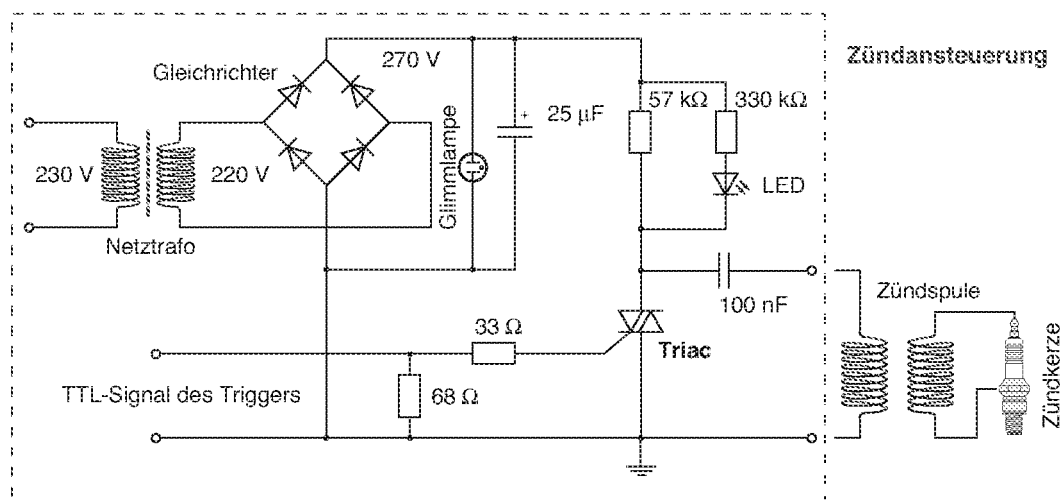


Abbildung 5.3: Die Zündansteuerung.

Ein Transformator liefert eine Wechselspannung von 220 V, die von einem *Brückengleichrichter* in eine pulsierende Gleichspannung umgewandelt und über einen *Elektrolytkondensator* (Elko) geglättet wird (siehe Kap. 5.2.1 und 5.2.2). Somit liegt an einem *Triac* (siehe Kap. 5.2.3) und an einem Kondensator mit der Kapazität von 100 nF eine Gleichspannung an. Sobald der Kondensator aufgeladen ist, fließen keine Ladungen mehr. Wenn ein Triggersignal auf das Gate des Triacs

gegeben wird, wird dieser leitend. Der Kondensator liegt dann auf Masse und kann sich über die Primärspule des Zündtransformators entladen. In der Sekundärspule wird eine Spannung induziert, die einen Funken an der Zündkerze auslöst.

Die beiden Widerstände vor dem Gate des Triacs sind *Abschlusswiderstände* (siehe Kap. 5.2.4), die üblicherweise $50\ \Omega$ betragen, hier aber so modifiziert wurden, dass der TTL-Pegel ausreicht, um den Triac durchzuschalten. Parallel zum Aufladewiderstand des Kondensators ist ein weiterer Widerstand mit einer LED in Reihe geschaltet, welche aufleuchtet, sobald ein Strom fließt. Sie dient also als optischer Hinweis auf einen Zündvorgang. Die Zündansteuerung befindet sich komplett auf einer Platine, die direkt an der Zündkerze angebracht ist, um einen möglichst geringen Abstand zwischen Zündspule und Zündkerze zu gewährleisten. Im Folgenden werden einige der elektronischen Bauteile genauer erläutert.

5.2.1 Der Gleichrichter

Wird für die Versorgung einer elektronischen Schaltung eine Gleichspannungsquelle benötigt, kann eine Versorgung aus dem Wechselspannungsnetz vorgenommen werden. Hierzu ist außer einer Transformation auch eine Gleichrichtung der Wechselspannung notwendig. Zur Realisierung werden Halbleiterdioden wegen ihrer Stromrichtungsempfindlichkeit, ihrer hohen Belastbarkeit und ihrer niedrigen Durchlassspannung verwendet. Sie bilden hier eine *Vollweggleichschaltung*, die beide Halbwellen der Eingangswechselspannung ausnutzt (siehe Abb. 5.4). Der sich anschließende Elektrolytkondensator (siehe Abb. 5.3 und Kap. 5.2.2) wird also in einer solchen Schaltung während einer Periode zweimal aufgeladen. Bei der positiven Halbwellen der Eingangsspannung U_e liegen die Dioden D_1 und D_4 in Durchlassrichtung, bei der negativen Halbwellen die Dioden D_2 und D_3 . Die Ausgangsspannung U_a entspricht deshalb der positiven und der invertierten negativen Halbwellen der Eingangsspannung, jedoch für beide Halbwellen vermindert um die doppelte Durchlassspannung der Dioden D_1, D_4 bzw. D_3, D_2 .

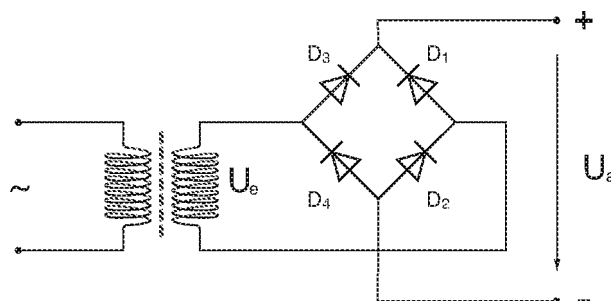


Abbildung 5.4: Der Brückengleichrichter mit Transformator.

5.2.2 Der Elektrolytkondensator

Elektrolytkondensatoren haben im Gegensatz zu anderen Kondensatoren eine sehr hohe Kapazität bei vergleichsweise kleiner Bauform. Die Funktionsweise ist prinzipiell die gleiche wie die eines Plattenkondensators, jedoch ist eine der Platten durch einen Elektrolyten, eine leitende Flüssigkeit, ersetzt. Die andere Platte wird durch eine Oxid-Schicht (Dielektrikum) isoliert. Der Elektrolytkondensator ist polungsabhängig: eine falsche Polung würde ihn zerstören.

5.2.3 Der Triac

Ein Triac ist ein Halbleiterbauelement, das in vielen Anwendungen mit erheblichen Vorteilen anstelle mechanischer Schalter eingesetzt werden kann. Die Bezeichnung „Triac“ ist eine Abkürzung für „**Tri**ode **alternating current switch**“. Das bedeutet, dass es sich bei diesem Bauelement um einen Wechselstromschalter mit drei Anschlüssen handelt: Das Gate dient zum Schalten des Triacs, die beiden anderen Elektroden werden als Anode 1 und Anode 2 bezeichnet (siehe Abb. 5.6). Man kann sich den Triac als eine Kombination zweier *Thyristoren* vorstellen. Ein Thyristor hat, ähnlich wie eine Diode, eine Durchlass- und eine Sperrrichtung, wird im Gegensatz zu ihr allerdings nur leitend, wenn an seine dritte Elektrode, auch Gate genannt, eine geeignete Spannung angelegt wird. Der Triac setzt sich aus zwei antiparallel geschalteten Thyristoren (siehe Abb. 5.5) zusammen, von denen einer am mittleren p-Leiter angesteuert wird, also bei positiven Steuerspannungen leitend wird. Der zweite Thyristor wird am mittleren n-Leiter angesteuert, schaltet also bei negativen Steuerspannungen durch. Die beiden Steuerelektroden sind miteinander verbunden, so dass zum Ansteuern eines Triacs nur ein Anschluss benötigt wird. Ein beliebig gepolter Impuls zwischen Gate und Anode 1 schaltet den Triac durch, unabhängig von der Polung der Spannung im Ladestromkreis.

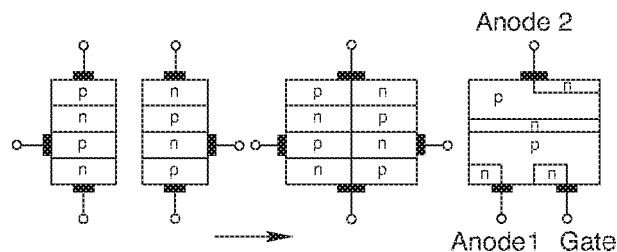


Abbildung 5.5: Aufbau eines Triacs aus zwei antiparallel geschalteten Thyristoren.

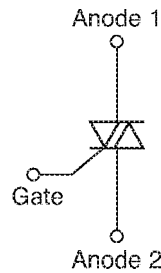


Abbildung 5.6: Schaltzeichen eines Triacs.

5.2.4 Abschlusswiderstand

Eine häufige Ursache für Fehlfunktionen in Hochfrequenz-Leitungen sind Leitungsreflexionen. Jede Leitung, unabhängig von ihrer Geometrie, hat einen induktiven, kapazitiven und ohmschen Anteil am Wechselstromwiderstand, wodurch Schwingungen im Signal auftreten können. Der ideale Übergang am Ende einer Übertragungsleitung lässt keinerlei Störung des Signals zu. In diesem Fall spricht man von einer abgeschlossenen Leitung und es gilt: $R_L = Z_0$ (Z_0 ist der Wellenwiderstand, R_L der Abschlusswiderstand). Ist die Leitung nicht mit dem Leitungswiderstand abgeschlossen ($R_L \neq Z_0$), so wird das Spannungssignal am Leitungsende reflektiert. Ein Maß für die Störung des Signals ist der so genannte Reflexionskoeffizient, der gerade für $R_L = Z_0$ null wird, d. h. wenn die Leitung sowohl am Eingang als auch am Ausgang angepasst ist (siehe dazu auch [HBG94]).

5.3 Aufbau der Zündschaltung und -ansteuerung

Die Konstruktionspläne für die Funkenkammer, die vom Europäischen Kernforschungszentrum CERN stammen, beinhalten einen Vorschlag für die Komponenten der Zündschaltung (wie Kondensatoren), die teilweise für den Bau dieser Funkenkammer verwendet wurden, ebenso wie die Spezifikationen des mit Stickstoff gefüllten Gehäuses für die Funkenstrecke. Eine Zündkerze wurde im Institut für Materialphysik der Universität Münster für unsere Zwecke modifiziert.

Die in Kap. 5.2 beschriebene Zündansteuerung wurde auf einer Platine angebracht, die direkt an der Zündkerze montiert ist (siehe Abb. 5.7). Das gewährleistet einen geringen Abstand zwischen der in der Mitte der Platine sitzenden Zündspule und der Zündkerze, die aus dem mit Stickstoff gefüllten Gehäuse ragt.

Die Zündschaltung wurde zunächst mit Hilfe eines Pulser getestet, der das Triggersignal simulierte, welches später auf die Zündansteuerung gegeben werden sollte. Mit dem Pulser wurde bei unterschiedlichen Hochspannungen ein Signal auf die

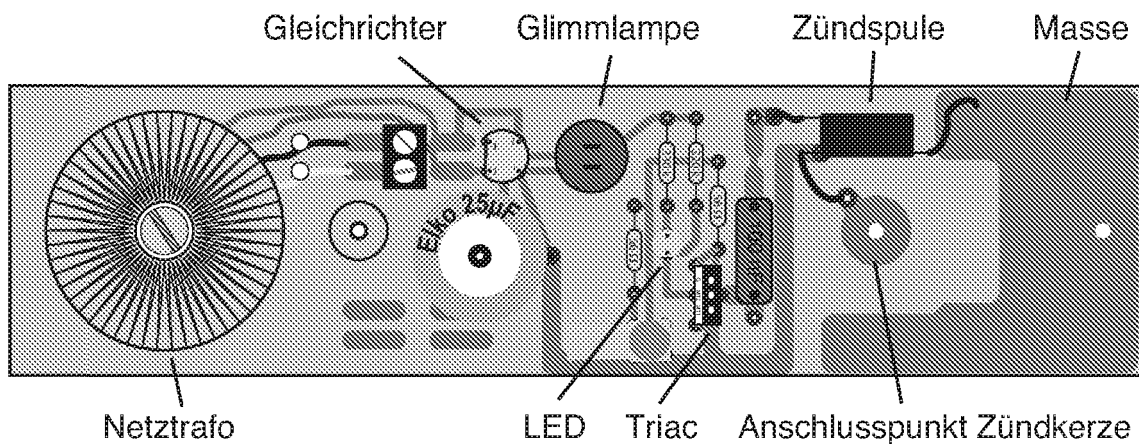


Abbildung 5.7: Die Zündansteuerung.

Zündansteuerung gegeben, wobei bis zu einer anliegenden Spannung von 2 kV lediglich ein leises Zündgeräusch an der Zündkerze zu registrieren war. Oberhalb dieser Spannung gab es Funkenüberschläge zwischen Zündkerze und Messingelektrode in dem Gehäuse und die Kondensatoren wurden entladen. Bei weiterem Erhöhen der Spannung wurde das Zündgeräusch immer lauter und das Gehäuse leuchtete seitlich hell auf. Wurde eine Hochspannung von 7 kV erreicht, so fanden selbstständige Entladungen ohne Betätigen des Pulsers statt. Der Grund dafür war, dass für diese Spannung die Zündkerze zu nah an der Elektrode lag und eine Entladung ohne die Hilfsfunkenstrecke stattfinden konnte. Der Abstand wurde daraufhin mittels einer Unterlegscheibe vergrößert, die um das Gewinde der Zündkerze gelegt wird, so dass der Kopf nicht mehr so weit in das Gehäuse ragt. Dies ermöglicht das Anlegen noch höherer Spannungen.

Während des Tests der Zündansteuerung mit dem Pulser waren natürlich noch keine Funken in der Kammer zu sehen, da der Zeitpunkt des Pulsersignals unabhängig von einem möglichen Teilchendurchgang ist und nur eine zufällige Koinzidenz zwischen Teilchendurchgang und Pulsersignal zu der Darstellung einer Bahnspur in Form von Funken hätte führen können.

Auch als das Triggersignal zur Ansteuerung benutzt wurde, gab es zunächst noch keine sichtbaren Entladungen. Um mögliche Ursachen hierfür ausfindig zu machen, wurde zunächst mit Hilfe eines Oszilloskops die Zeitverzögerung zwischen Teilchendurchgang und Anlegen der Hochspannung bestimmt, der Abstand zwischen Zündkerze und Elektrode variiert und die Einstellungen der Elektronik (Hochspannung der Photomultiplier, Schwellen der Diskriminatoren) überprüft.

Einen großen Einfluss auf den zeitlichen Abstand zwischen Teilchendurchgang und Anlegen der Hochspannung hat die Induktivität der Zündspule. Dazu wurde

eine Reihe von Tests mit unterschiedlichsten Zündspulen vorgenommen. Die Geschwindigkeit der Strominduktion wurde für mehrere handelsübliche Spulen wie zum Beispiel eine Standard-Autozündspule gemessen und auch für selbstgewickelte Spulen mit, und zur Erhöhung der Strominduktionsgeschwindigkeit auch ohne, Eisenkern. Letztendlich lieferte die Zündspule für eine Blitzlampe das beste Ergebnis (siehe Anh. A). Mit ihr liegt die Gesamtverzögerung zwischen Triggersignal und Hochspannungsimpuls bei etwas unter $1\ \mu\text{s}$.

Nach dem Aufbau der Zündansteuerung bedurfte es einer längeren Spülung mit Neon und Helium sowie eines erneuten Abdichtens der Kammer, bis erste einzelne Funkenüberschläge und mit der Zeit immer deutlicher ausgebildete Funkenabfolgen zu beobachten waren.

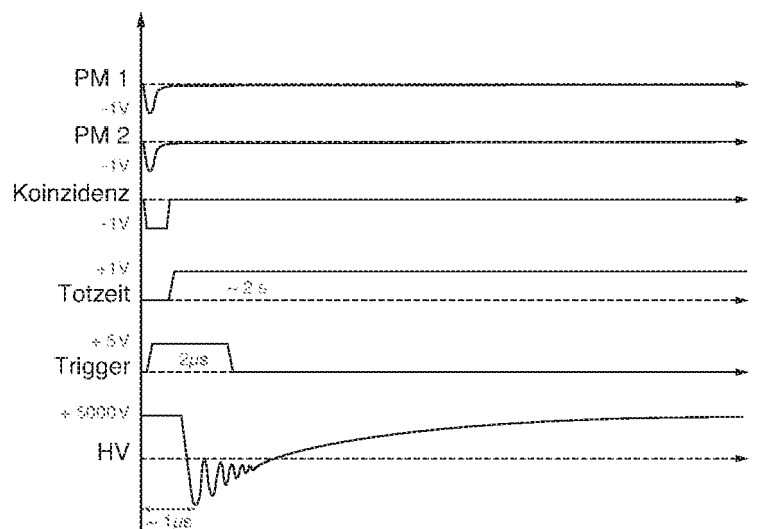


Abbildung 5.8: Das zeitliche Verhalten vom Teilchendurchgang zur Funkenentladung (schematisch).

Das in Abb. 5.8 dargestellte Zeitdiagramm zeigt schematisch den zeitlichen Verlauf der Signale zwischen Teilchendurchgang und Anlegen der Hochspannung an den Platten der Funkenkammer. Die Verzögerung der Signale durch die NIM-Module, wie Diskriminatoren, Koinzidenzeinheit, Timer und NIM-TTL-Wandler, beträgt jeweils einige Nanosekunden. Zusammen mit der Verzögerung zwischen Eingehen des Triggersignals (aus dem NIM-TTL-Wandler) und dem Hochspannungsimpuls erhält man die Gesamtverzögerung von etwa $1\ \mu\text{s}$.

5.4 Optimierung der Funkenkammer für den Dauerbetrieb

Wird die Funkenkammer bei der Inbetriebnahme mit Gas (75% Neon und 25% Helium, vergleiche Kap. 3.1) gespült, ist eine stetige Verbesserung der Funktionsweise erkennbar. Anfangs waren ausschließlich Funkenüberschläge an den gleichen Stellen in der Kammer zu beobachten. Diese Entladungen waren unabhängig von Bahnspuren ionisierender Teilchen. Solche Vorzugsstellen können *Spitzen* (kleine Unregelmäßigkeiten) in der Plattenkaschierung sein, an der das elektrische Feld zwischen den Platten nicht homogen ist. Ist eine lokale Erhöhung des elektrischen Feldes die Folge, so ist die Wahrscheinlichkeit für eine Funkenentladung hier sehr groß.

Es fiel auf, dass die Überschläge vorzugsweise an den Rändern der Platten stattfanden sowie an den Kontaktstellen der Platten, wo Anschlussstecker angreifen, um Ladung auf die Platten zu leiten. An diesen Stellen ist der Abstand zwischen den Elektroden geringer als an anderen. Auch nach ausreichender Gasspülung erkannte man nach der ersten Inbetriebnahme immer wieder überdurchschnittlich viele Entladungen an diesen Vorzugsstellen. Um dieses Verhalten zu reduzieren, wurden nach der ersten Inbetriebnahme erneut alle Platten aus der Kammer entfernt und in folgender Weise nachbearbeitet:

Sie wurden zunächst mit Stahlwolle und anschließend mit sehr feinkörnigem Schleifpapier abgeschliffen, wobei die Ränder der Platten abgerundet wurden. Auf Höhe der Anschlussstecker wurde an den jeweils gegenüberliegenden Platten ein quadratisches Stück der Fläche $17 \times 17 \text{ mm}^2$ der Kupferlegierung ausgefräst, um den Abstand der gegenüberliegenden Elektroden zu diesen Steckern zu vergrößern. Um nicht zusätzliche scharfe Kanten und somit neue Feldinhomogenitäten zu erzeugen, wurden die Ecken dieser ausgefrästen Stücke anschließend abgerundet. Auch die Ränder der Platten mussten sorgfältig glatt geschliffen werden. Anschließend wurden sie gereinigt und mit einer Metallpolitur bearbeitet. Eine erneute Inbetriebnahme zeigte eine starke Reduzierung der unerwünschten Entladungen.

Zusammenfassung

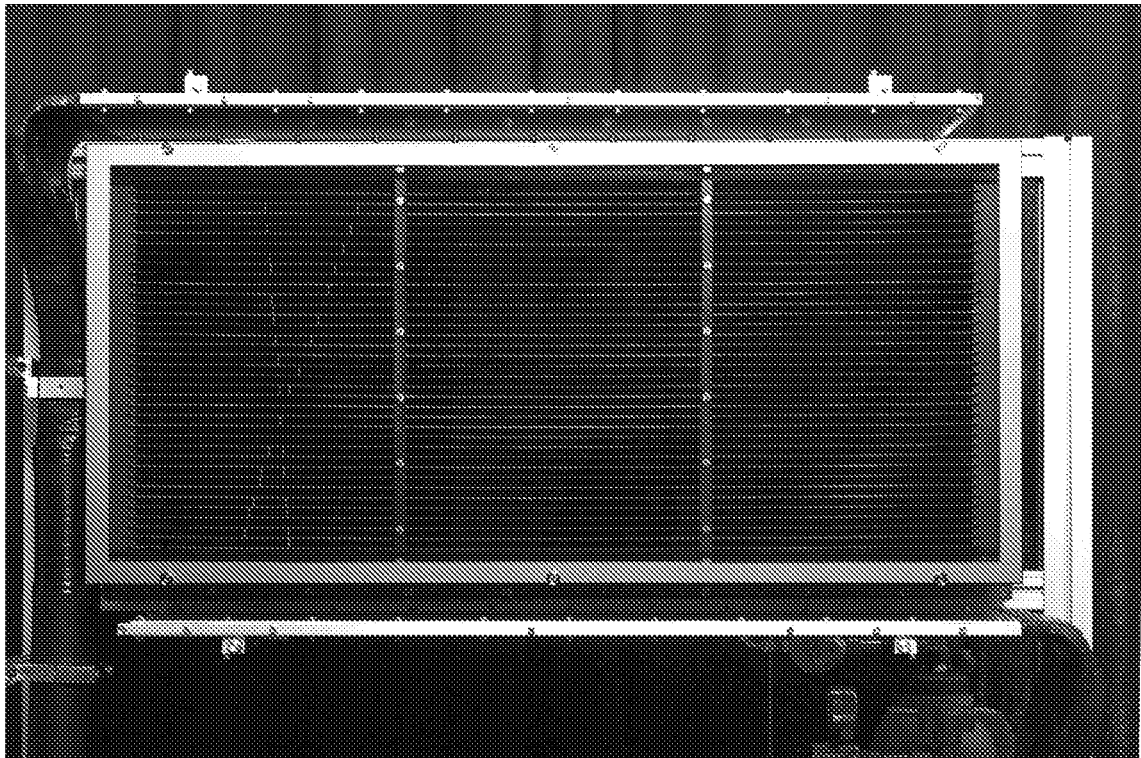


Abbildung 5.9: Fotoaufnahme der Funkenkammer.

Im Rahmen der vorliegenden Arbeit wurde eine Funkenkammer aufgebaut und in Betrieb genommen. Mit diesem Detektor können hochenergetische Teilchen der sekundär-kosmischen Strahlung in Form von Funkensequenzen nachgewiesen werden.

Die Konstruktionspläne für die Funkenkammer wurden am Europäischen Kernforschungszentrum CERN entwickelt. In Eigenarbeit wurden das Triggersystem inklusive der Triggerdetektoren und der Elektronik entwickelt und gebaut, sowie die Zündansteuerung, die Gasversorgung und Komponenten der Zündschaltung.

Das Gestell, ein so genanntes Kanya-Verbindungssystem aus Aluminium, wurde so konzipiert, dass die Funkenkammer daran befestigt werden konnte und im unteren Teil die Elektronik sowie das Gassystem inklusive der beiden Gasflaschen Platz finden. Es steht auf Rollen, so dass die Kammer mobil ist.

Bei der Inbetriebnahme der Funkenkammer traten zunächst unterschiedlichste Schwierigkeiten auf, die überwunden werden mussten:

Als die Funkenkammer nach dem Aufbau des Triggers, der Zündschaltung, der Elektronik und des Gassystems erstmalig in Betrieb genommen werden sollte, waren zunächst noch keine Funken zu beobachten. Erst ein erneutes Abdichten der Kammer brachte erste Erfolge.

Sowohl die Photomultiplier als auch die Paddles wurden durch Probeaufbauten getestet, doch nach der Anfertigung und Anbringung der Triggerdetektoren an der Funkenkammer war die zunächst gemessene Koinzidenzrate von 0,1 Hz wesentlich niedriger als die erwartete. Diese Koinzidenzrate war auch gleichzeitig die maximale Rate für Funkenentladungen in der Kammer. Erst nach einmaligem Erhöhen der am oberen Photomultiplier anliegenden Spannung kam es zu einer Verbesserung der Koinzidenzrate auf 3,6 Hz, die jetzt stabil gemessen wird. Dies ist ein Indiz für Probleme mit der Base des oberen Photomultipliers, die auf Dauer ersetzt werden sollte. Darüber hinaus ist die Lichtisolierung der Paddles durch Mylarfolie und Klebeband sehr empfindlich gegen äußere Einwirkungen. Die Paddles sind sorgsam zu behandeln und ihre Lichtdichtigkeit sollte regelmäßig überprüft werden.

Die erste Zündansteuerung der Funkenkammer war zu langsam und so wurde in mehreren Versuchen eine neue entwickelt, die die Hilfsfunkenstrecke für die Zündschaltung ausreichend schnell auslöst. Entscheidend war die Wahl der Zündspule sowie der räumliche Abstand zwischen ihr und der Zündkerze, um die Hochspannung in kürzester Zeit an die Platten der Funkenkammer zu legen.

Die Hochspannung, die im Falle eines Teilchendurchgangs an den Platten der Funkenkammer anliegen muss, um Funken zu erzeugen, ist von der Reinheit des Gasgemisches abhängig und beträgt nach ausreichendem Durchspülen des Kammergefäßes ca. 4,6 kV.

Die Funkenentladungen finden nun durch einen Timer begrenzt mit einer Rate von 0,5 Hz statt. Dabei sind neben eindrucksvollen Bahnspuren durch die ganze Kammer auch gelegentlich Teilchenschauer in der Kammer zu beobachten. Das Foto, das eine Abfolge von Funken in der Funkenkammer zeigt, wurde bei einer Belichtungszeit von 5 s aufgenommen (siehe Abb. 5.9). In dieser Zeit kam es zu zwei Zündvorgängen, die jeweils eine Funkenentladung zur Folge hatten.

A. Die technischen Daten der Funkenkammer

Die Kammer

Anzahl der Platten	37
Plattendicke	1,6 mm
Abstand der Platten zueinander	8 mm
Plattenlänge	80 cm
Plattenbreite	20 cm
Plattenmaterial	FR 4 Epoxy-Glasfaser-Laminat beidseitig mit Kupfer kaschiert

Die Photomultiplier

Typ	RCA 8575
Durchmesser	50,8 mm
Anzahl der Dynoden	12
Fenstermaterial	Pyrex-Glas (Typ 7740)
Dynoden	
Trägermaterial	CuBe
Sekundärelektronen emittierende Oberfläche	BeO
Anodenpuls Anstiegszeit, 2000 V	2,8 ns
Elektronen Durchgangszeit, 2000 V	37 ns
Maximale Hochspannung	3000 V
Absorptionswellenlänge	260 nm - 660 nm
λ (maximales Ansprechverhalten)	385 nm

Die Paddles

Szintillatormaterial	BC-408 (Bicron)
Szintillatorbreite	20 cm
Szintillatorlänge	72 cm
λ (Emissionsmaximum)	425 nm
Anstiegszeit des Pulses	0,9 ns
Abklingzeit des Pulses	2,1 ns
Abstand der Szintillatoren zueinander	50 cm
Lichtleitermaterial	Plexiglas (PMMA)
Lichtleiterfläche an der Szintillatorseite	$20 \times 1 \text{ cm}^2$
Lichtleiterfläche an der Photomultiplierseite	$3,56 \times 3,56 \text{ cm}^2$

Die NIM-Module

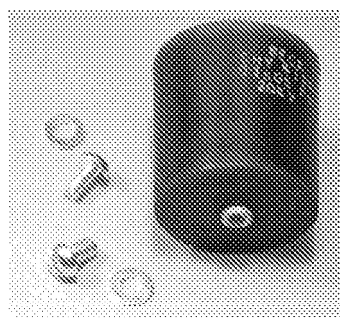
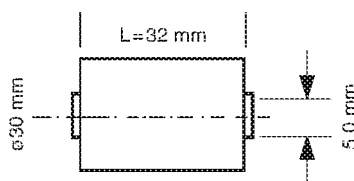
Hochspannungsquelle (Kammer)	C.A.E.N.	Mod. N 471
Hochspannungsquellen (PM)	Canberra	Mod. 3002
Dual Timer	C.A.E.N.	Mod. N 96B
NIM-TTL-Wandler	C.A.E.N.	Mod. 89
Diskriminatoren	GSI	CF 8000
Koinzidenzeinheit	GSI	CO 4000

Die Zündspule

Typ	ZS 1052-11 (Conrad Electronic GmbH)
Maximale Primärspannung	300 V
Ausgangsspannung sekundär	11 kV

Die 18 Kondensatoren zum Anlegen der Hochspannung

Typ	Hochspannungs-Plattenkondensator Keramik
	Best.-Nr. 10.11.81.910.5 (CERN)
Kapazität	1 nF
Maximale Spannung	20 kV

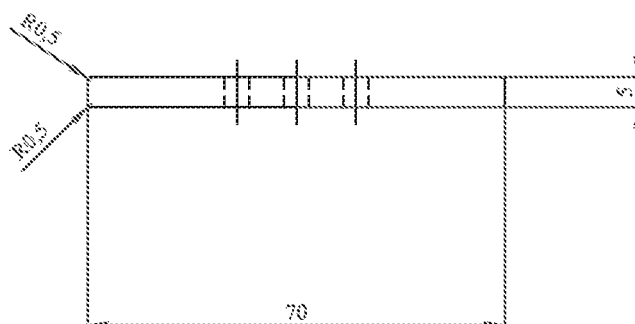
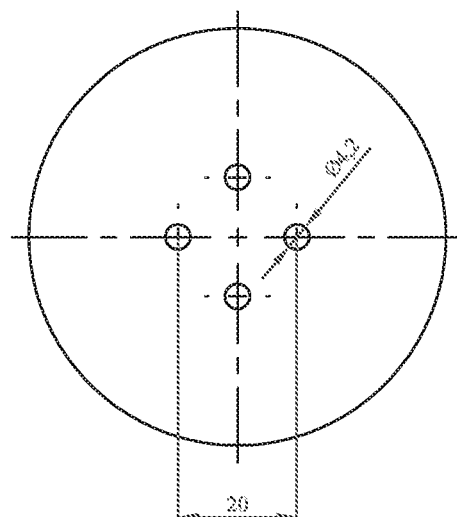


Einstellungen für den Dauerbetrieb

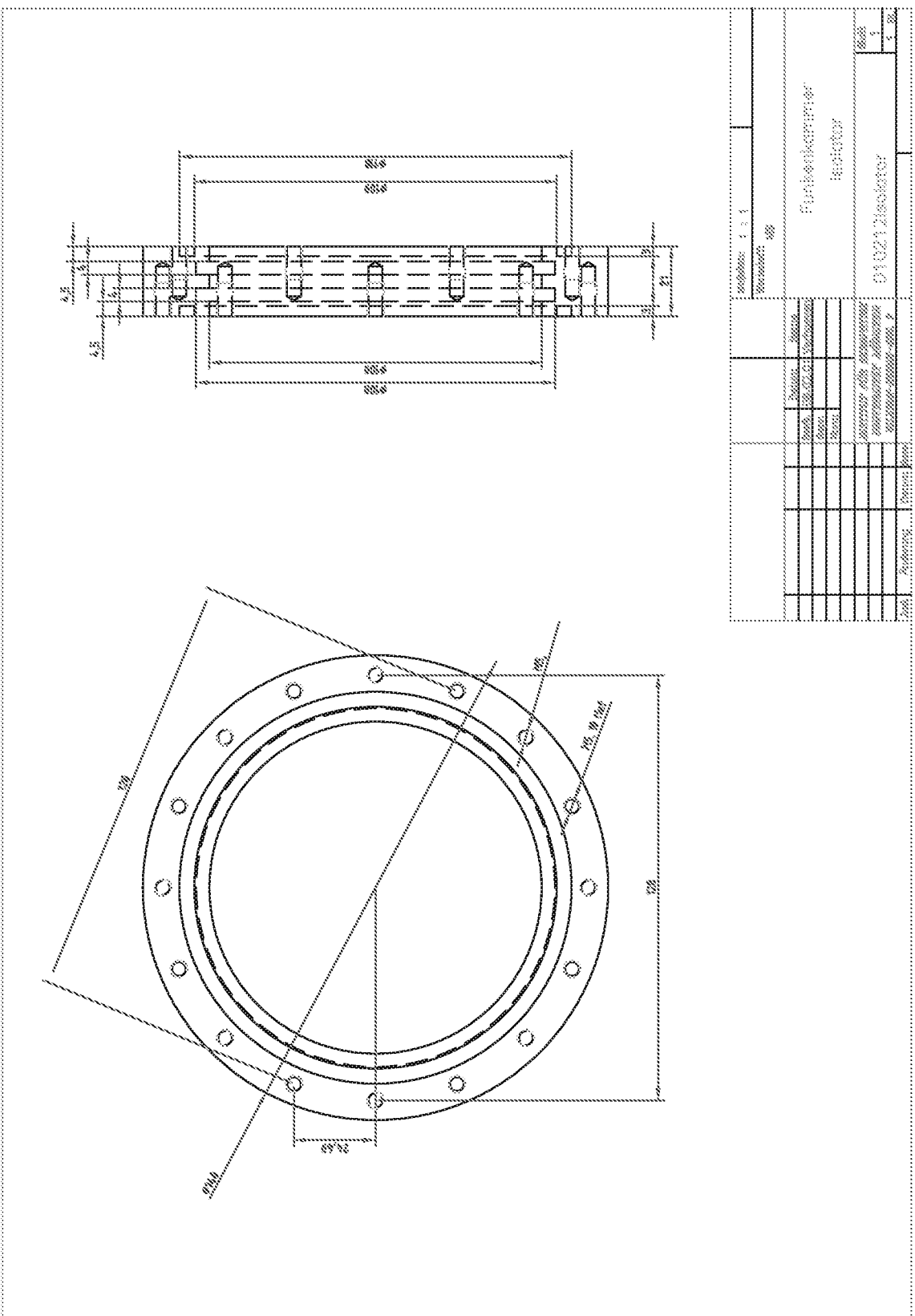
Arbeitsspannung am oberen Photomultiplier	2317 V
Arbeitsspannung am unteren Photomultiplier	2087 V
Arbeitsspannung an den Platten der Kammer	4600 V
Diskriminatorschwelle für den oberen Photomultiplier	107,6 mV
Diskriminatorschwelle für den unteren Photomultiplier	111,3 mV
Totzeit durch Timer	2 s

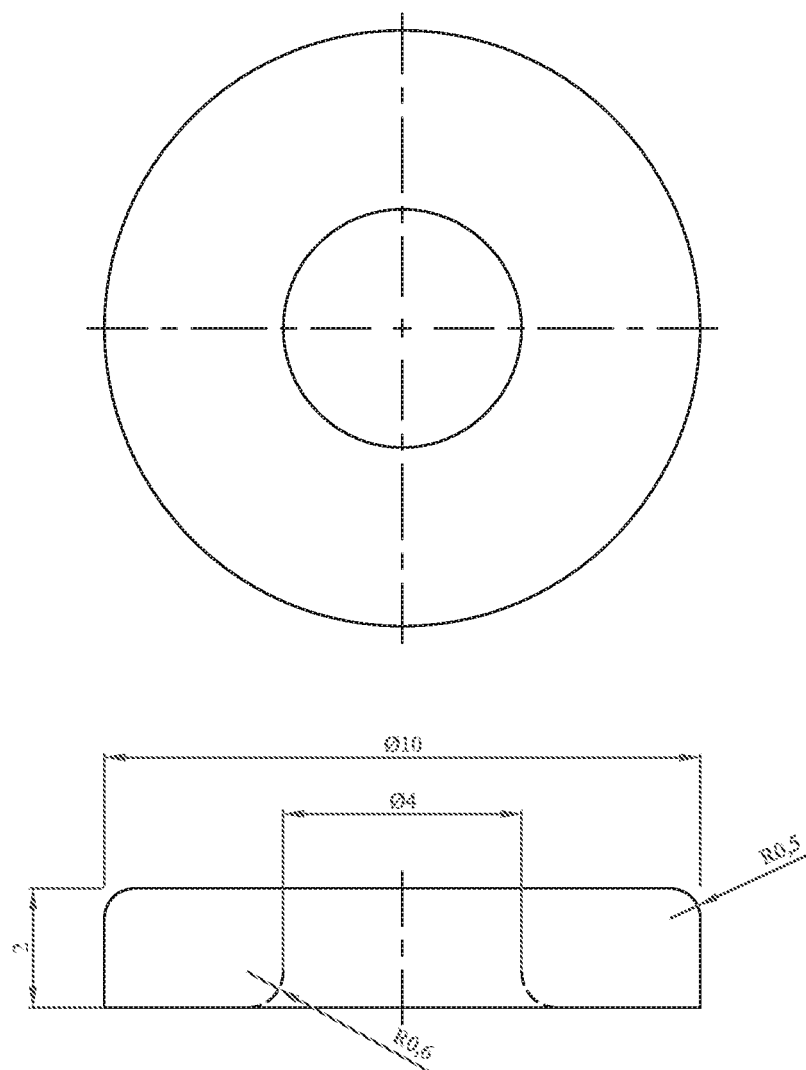
B. Die Konstruktionspläne

Im Folgenden sind die im Institut für Kernphysik der Universität Münster modifizierten bzw. entwickelten Konstruktionspläne der Funkenkammer aufgelistet.

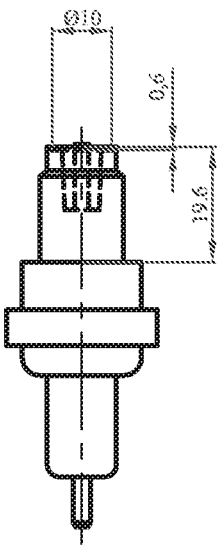


				Maßstab:		
				Verkauf:		MS
				Datum:	9.02.01	Verfahren
				Gepr.:		
				Norm:		
				INSTITUT FÜR KRAFTFORSCHUNG		
				UNIVERSITÄT MÜNSTER		
				WILHELM-KLAUW-STR. 8		
				010209Elektrode		Blatt 1
Zust.	Änderung	Datum	Notiz			1. Aufl.

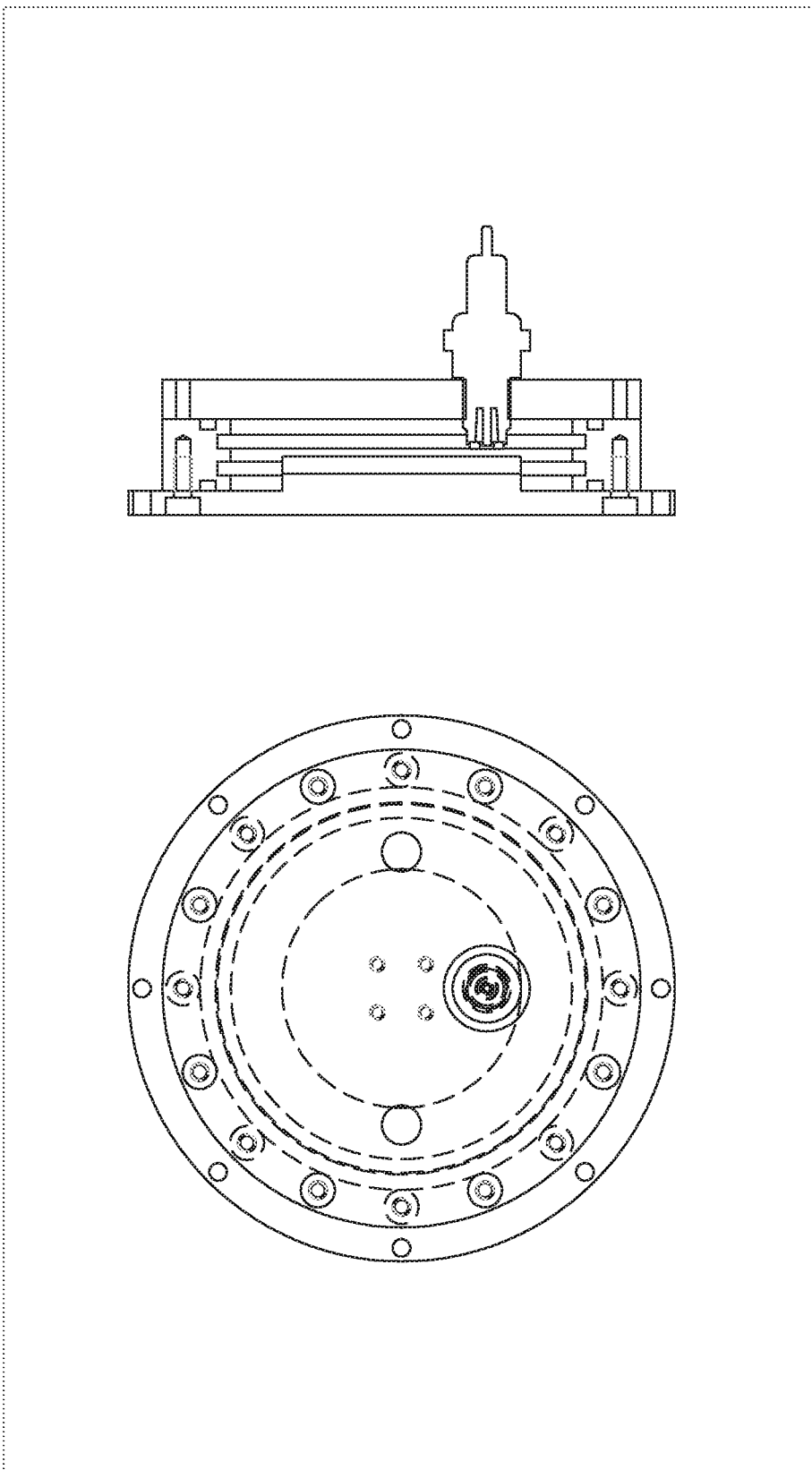
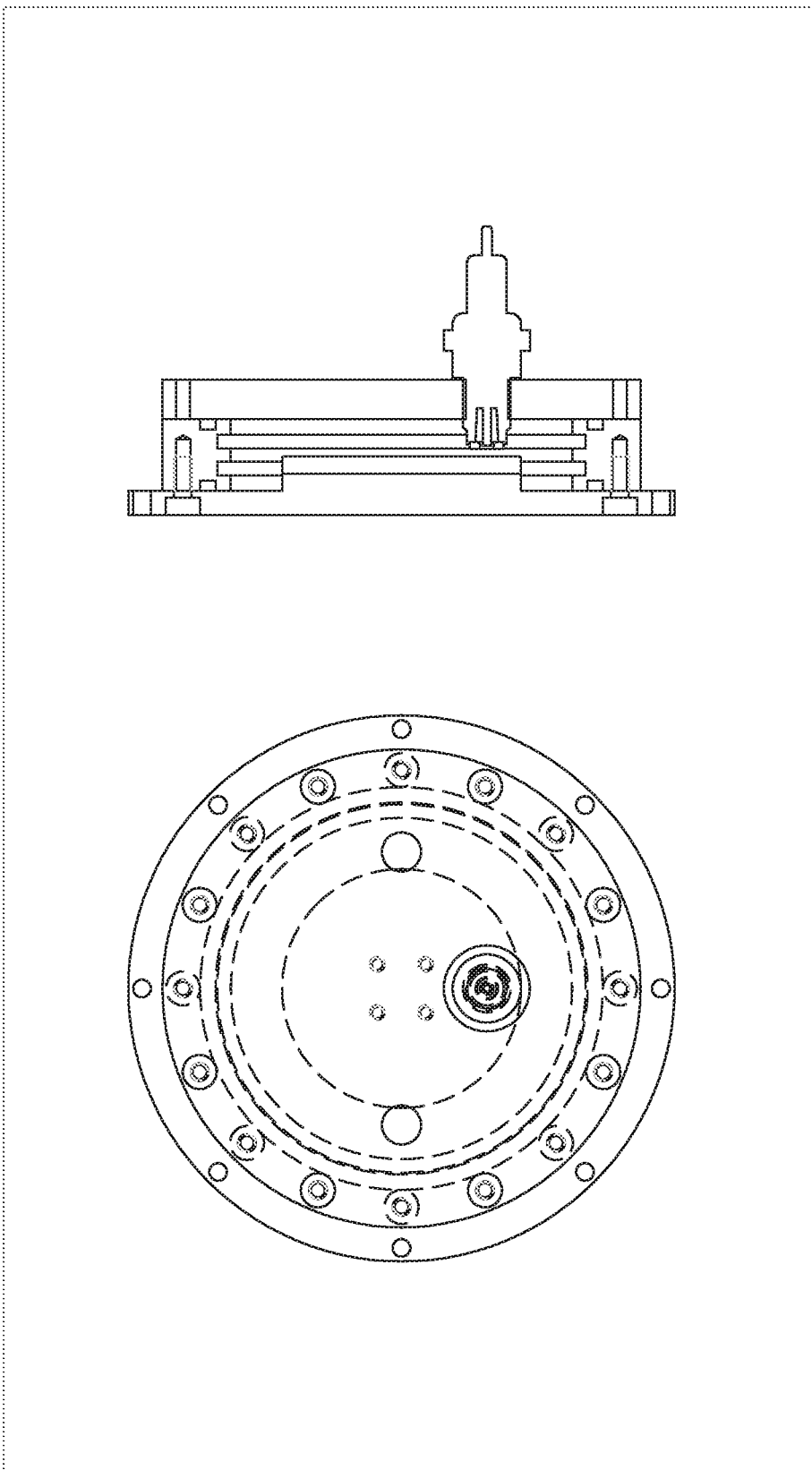




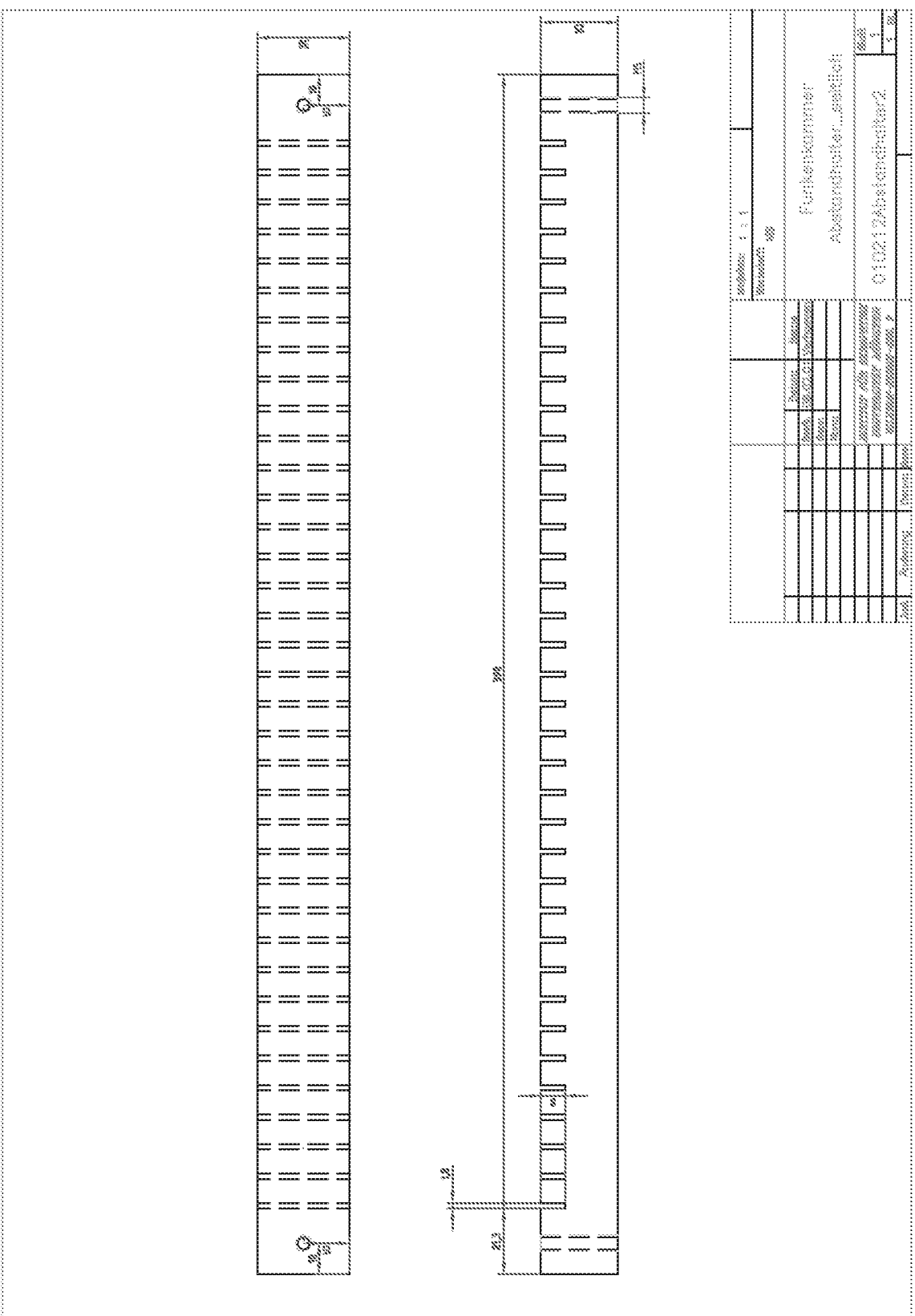
				Maßstab: 10:1	
				Werkstoff: Wolfram	
				Datum	Name
				Besch. 12.02.01	Verhoeven
				Gepr.	
				Name	
				INSTITUT FÜR KRAFTFORSCHUNG UNIVERSITÄT MÜNSTER WILHELM-KLAUW-STR. 8	
				010212Wolframelektrode	
				Blatt 1	
				1 von 1	
Zust.	Änderung	Datum	Name		

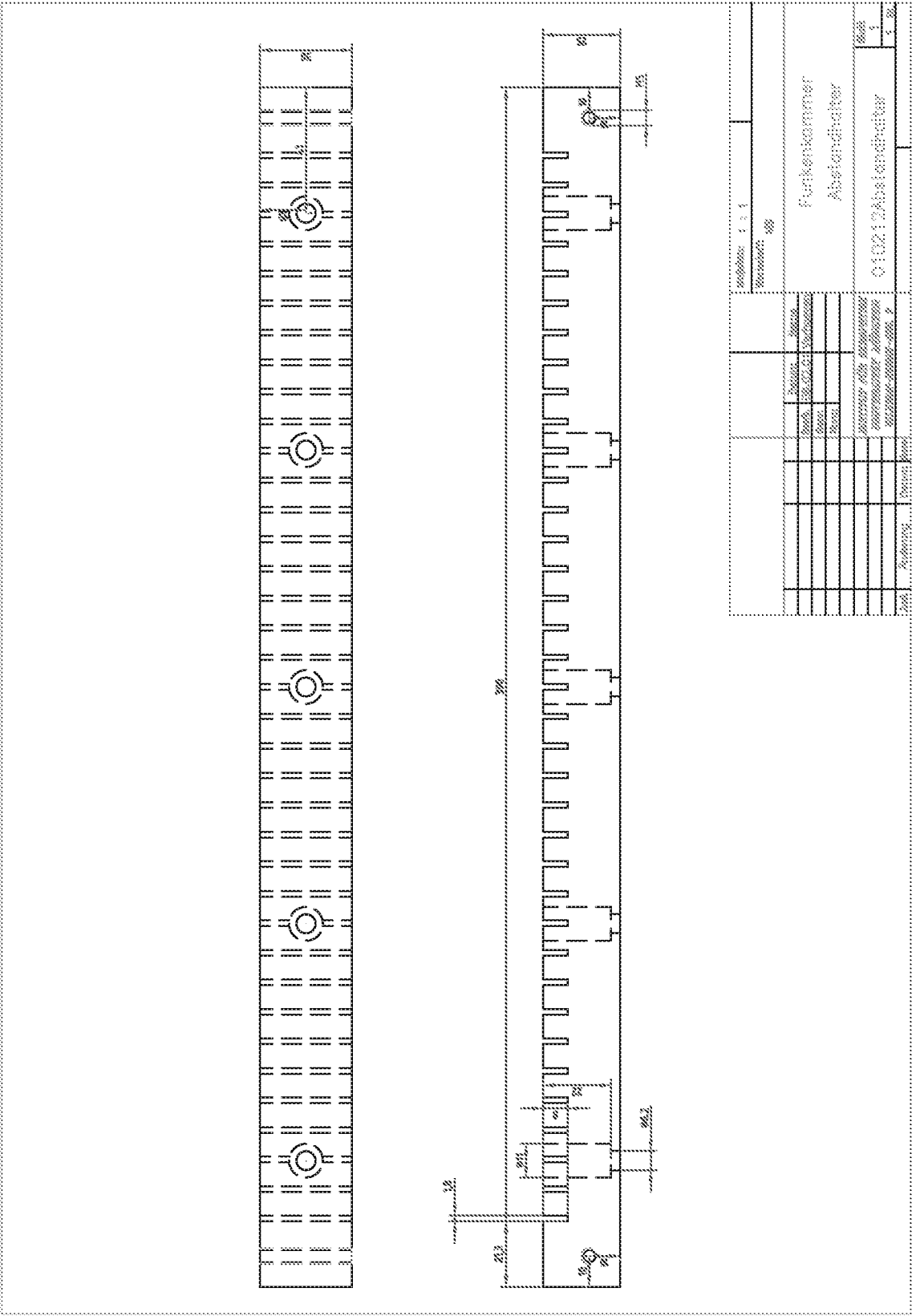


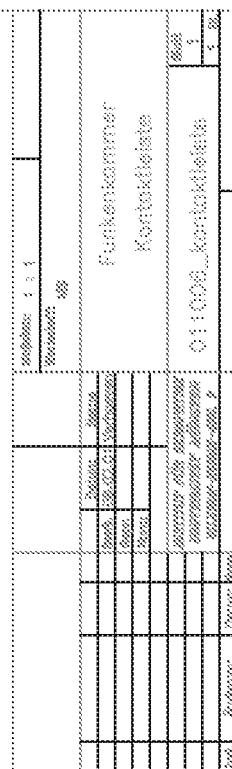
				Maßstab:			
				Werkstoff:			
				Zündkerze Bosch WRCC			
				Datum	Gezeichnet	Funkenkammer Zündkerze, modifiziert	
			Besch.	12.02.01	Verfahrenen		
			Gepr.				
			Norm				
						010212Zündkerze	
				INSTITUT FÜR KERNPHYSIK			
				UNIVERSITÄT MÜNSTER			
				KLEIN-KLEIN-STR. 8			
Zeit	Änderung	Datum	Name			Blatt 1 1 III	



The image contains two technical drawings of a mechanical assembly. The top drawing is a side view showing a central vertical shaft with a flange, mounted on a base with two vertical supports. The bottom drawing is a top view showing a circular flange with 12 bolt holes around its perimeter and a central circular feature with four small holes.







Literaturverzeichnis

- [All69] O. C. Allkofer, *Spark Chambers*, Thiemig-Verlag, 1969.
- [And36] C. D. Anderson, *Nobel Lectures Physics*, 365-376, World Scientific, 1998.
- [Dem95] W. Demtröder, *Experimentalphysik 2: Elektrizität und Optik*, Springer Verlag, 1995.
- [Eva55] R. D. Evans, *The Atomic Nucleus*, McGraw-Hill, New York, 1955.
- [Far65] F. J. M. Farley, *Progress in Nuclear Techniques and Instrumentation 1*, North-Holland, 1965.
- [Fis96] G. J. Fishman, *IAU Symp. 165: Compact Stars in Binaries*, 467, 1996.
- [Ger93] C. Gerthsen, H. Vogel, *Physik*, Springer Verlag, 1993.
- [Gre66] K. Greisen, *Phys. Rev. Lett.* **16** (1966).
- [Gru85] C. Grupen, *Kosmische Strahlung*, Physik in unserer Zeit **3**, 1985.
- [Gru96] C. Grupen, *Particle Detectors*, Cambridge 1996.
- [Hag02] K. Hagiwara *et al.*, *Phys. Rev. D* **66**, 010001 (2004).
- [HBG94] E. Hering, K. Bressler, J. Gutekunst, *Elektronik für Ingenieure*, VDI Verlag, 1994.
- [Hes36] V. F. Hess, *Nobel Lectures Physics 2*, 360-362, World Scientific, 1998.
- [Kle87] K. Kleinknecht, *Detektoren für Teilchenstrahlung*, Teubner Verlag, 1992.
- [Kub04] V. Kubarovsky *et al.*, *Phys. Rev. Lett.* **92** (2004).
- [Leo87] W. R. Leo, *Techniques for Nuclear and Particle Physics Experiments*, Springer Verlag, 1992.

- [Lon92] M. S. Longair, *High Energy Astrophysics I+II*, Cambridge University Press, Cambridge, 1992 & 1994.
- [Mor96] C.E. Mortimer, *Chemie: Das Basiswissen der Chemie*, Georg Thieme Verlag, 1996.
- [Mül03] M. Müller, *Untersuchung unbegleiteter Hadronen in durch kosmische Strahlung induzierten Luftschauern im Bereich bis zu einem PeV*, Wissenschaftliche Berichte, FZKA 6912, 2003.
- [Ori02] Auszüge aus Handbuch nach D. Orie, O E Technologies, 2002.
- [PDG00] D.E. Groom *et al.* (Particle Data Group Collaboration), *Review Of Particle Physics*, Eur. Phys. J. C **15** (2000).
- [Per90] D.H. Perkins, *Introduction to High Energy Physics*, Addison-Wesley, 1987.
- [PRS99] B. Povh, K. Rith, C. Scholz, F. Zetschke, *Teilchen und Kerne*, Springer Verlag, 1999.
- [Rae63] H. Raether, *Proceedings of the International Conference on Ionization Phenomena in Gases*, Paris, 1963.
- [Ric74] P. Rice-Evans, *Spark, Streamer, Proportional and Drift Chambers*, London, 1974.
- [Roe96] B.P. Roe, *Particle Physics at the New Millenium*, Springer Verlag, 1996.
- [Rut11] E. Rutherford, *The Scattering of α and β Particles by Matter and the Structure of the Atom*, Philosophical Magazine **21** 669 (1911).
- [Rut64] J.G. Rutherglen, Progr. Nucl. Phys. **9** (1964).
- [Wu57] S. Wu, E. Ambler, R.W. Hayward, D.D. Hoppes, R.F. Hudson, Phys. Rev. **105** (1957).

Danksagung

An dieser Stelle möchte ich mich bei allen bedanken, die zum Gelingen meiner Diplomarbeit beigetragen haben.

Herrn Prof. Dr. Rainer Santo danke ich für die interessante und abwechslungsreiche Aufgabenstellung, die hervorragenden Arbeitsbedingungen am Institut für Kernphysik sowie die Ermöglichung der Teilnahme an der DPG-Tagung in Tübingen.

Bei Dr. Damian Bucher möchte ich mich für die Betreuung der Arbeit, seine hilfreichen Anregungen und stete Diskussionsbereitschaft bedanken.

Ein großes Dankeschön geht an Christian Klein-Bösing, Dr. Klaus Reygers und Prof. Dr. Johannes P. Wessels, die mit ihren Ideen und Diskussionen einen großen Anteil am Gelingen dieser Arbeit haben.

Ein besonderes Dankeschön möchte ich an Norbert Heine richten, nicht nur für die professionelle Unterstützung bei der Gestaltung der Bilder, der Diagramme und der Poster, für das Begleiten fast aller Arbeitsschritte wie das Formen der Paddles, den Aufbau der Zündansteuerung, das Ausfräsen der Platten und vieles mehr, sondern auch für viele wertvolle Ratschläge elektronischer und musikalischer Art.

Für seinen Einsatz bei der Anfertigung des Gassystems und die gute Zusammenarbeit bedanke ich mich bei Wolfgang Verhoeven.

Dank auch an Dr. Karl-Walter Glitza und Bernd-Holger Witt von der Universität Wuppertal für ihre Hilfsbereitschaft.

Ebenso möchte ich mich bei H.-H. Adam, Dr. S. Bathe, C. Baumann, Dr. H. Büsching, Dr. R. Glasow, H. Gottschlag, PD Dr. A. Khoukaz, T. Korfsmeier, N. Lang, R. Menke, T. Mersmann, T. Rausmann, B. Sahlmüller, S. Steltenkamp, A. Täschner, A. Wilk, O. Zaudtke und den Mitarbeitern der E-Werkstatt für die angenehme Zusammenarbeit und die Unterstützung bedanken.

Für die kritische Durchsicht der Arbeit danke ich nochmals Dr. Damian Bucher, Christian Klein-Bösing und Dr. Klaus Reygers sowie Dr. André Wenning und Kathrin Schröer. Ein herzliches Dankeschön an Vera Tendahl für die Hilfestellung bezüglich der neuen Rechtschreibung.

Ein Dankeschön geht an meine Freunde und Studienkollegen, die mich auch außerhalb der Physik begleitet haben: Alex, André, Anja, Boris, Christian, Elli, Floh, Isa,

Lars, Mareike, Nicola, Nicole, Ricki, Timo, Tobias A., Tobias W., Vera und Zoltan. Darüber hinaus möchte ich mich an dieser Stelle bei Frau Schenk für ihre Großzügigkeit bedanken.

Ein ganz besonderer Dank geht an Christian für seine Unterstützung nicht nur in physikalischen Belangen.

Abschließend danke ich meinen Großeltern, Oma Fine, Oma Sefi und beiden Opa Willis sowie meinem Vater, die mich in moralischer und oftmals auch finanzieller Hinsicht unterstützt haben.

Mareike, Männi und Petra haben mir in schwierigen Zeiten immer wieder Mut gemacht, dafür ein großes Dankeschön!

*Hiermit bestätige ich, dass ich diese Arbeit selbstständig
verfasst und keine anderen als die angegebenen
Quellen und Hilfsmittel verwendet habe.*

Münster, 3. März 2004

Melanie Hoppe