

EUROPEAN ORGANIZATION FOR NUCLEAR RESEARCH

NP Internal Report 69-24
28. Juli 1969

DIPLOMARBEIT

von

Robert Klanner

BAU UND BETRIEB VON WIDE-GAP DRAHTFUNKENKAMMERN MIT

MAGNETOSTRIKTIVEM AUSLESESYSTEM

<u>INHALTSVERZEICHNIS</u>		<u>Seite</u>
I.	EINLEITUNG	1
II.	FUNKENKAMMERN	1
	1) Einleitung, Prinzip und Definition der wichtigsten Parameter	1
	2) Theorie der Gasentladung	4
	3) Die verschiedenen Funkenkammern	7
	a) Streamer Kammern	7
	b) Kammern bei denen die Funkenden Teilchenbahnen folgen	10
	α) Funkenkammern mit kleinem Elektrodenabstand	11
	β) Funkenkammern mit grossem Elektrodenabstand	12
	4) Die gebauten Kammern	14
	5) Die Spannungsversorgung	17
III.	AUSLESESYSTEME VON FUNKENKAMMERN	19
	1) Die verschiedenen Systeme	19
	a) Photographische Methoden	19
	b) Digitalisierte Methoden	19
	α) Auslesesysteme für Kammern mit Plattenelektroden	19
	β) Auslesesysteme für Kammern deren Elektroden aus Drahtebenen gebildet sind	20
	2) Das magnetostriktive Signal	22
	3) Das verwendete Auslesesystem	26
IV.	MESSUNGEN UND TESTRESULTATE	34
	1) Rekonstruktion der Spuren	34
	2) Nachweiswahrscheinlichkeit	36
	3) Gedächtniszeit	38
V.	ANWENDUNG DER FUNKENKAMMERN IM CERN-BOSONSPEKTROMETER	40
	1) Die Methode der fehlenden Masse	40
	2) Experimenteller Aufbau, Messmethode und Abschätzung der Messfehler	43
	3) Messresultate	48

4) Vergleich der Messresultate mit Ergebnissen anderer Experimente und mögliche Interpretationen der gemessenen Strukturen	51
--	----

LITERATURANGABEN

59

I. EINLEITUNG

Diese Arbeit ist - im Rahmen einer Diplomarbeit an der Technischen Hochschule München in Experimentalphysik - das Ergebnis einer einjährigen Mitarbeit des Autors in der CERN-Bosonspektrometergruppe. Die Tätigkeit des Autors beschränkte sich dabei im Wesentlichen auf den Bau und den Betrieb von Funkenkammern und deren Auslesesystem, deren Beschreibung auch den Schwerpunkt der Arbeit bildet.

Das zweite Kapitel handelt ausschliesslich von Funkenkammern: Nach der Definition der wichtigsten Parameter wird kurz auf die Theorie der Gasentladung eingegangen. Mit ihrer Hilfe werden die Unterschiede zwischen den einzelnen Funkenkammertypen diskutiert. Anschliessend werden die gebauten Kammern genauer beschrieben.

Das dritte Kapitel beschäftigt sich mit dem Auslesen der Information aus Funkenkammern. Nach einer kurzen Diskussion der verschiedenen üblichen Systeme wird das magnetostriktive etwas genauer behandelt. Das nächste Kapitel berichtet über die Eigenschaften der in der Gruppe gebauten Kammern.

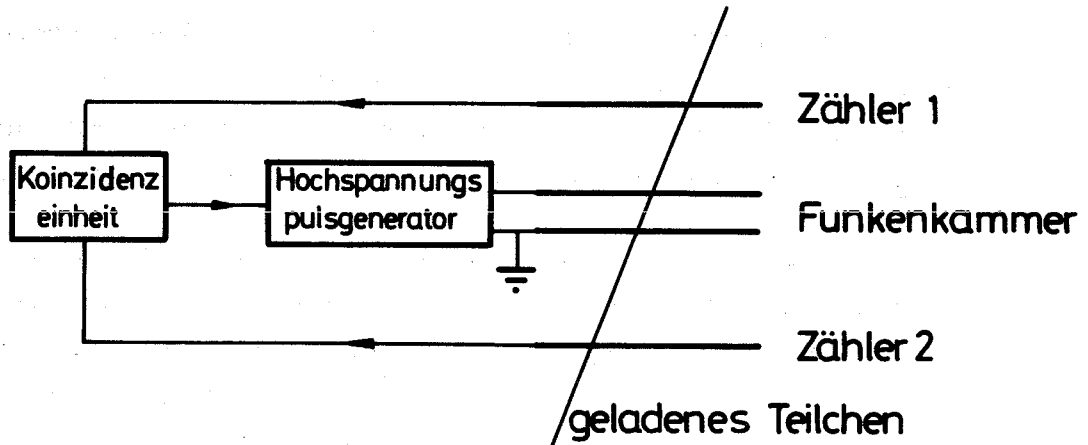
Das letzte Kapitel beschreibt das CERN-Bosonspektrometer, bei dem die gebauten Funkenkammern eingesetzt wurden. Nach einer Erläuterung der Messmethode und Abschätzung der Messgenauigkeit werden die im Jahre 1968 gemessenen Daten dargelegt und mögliche Interpretationen der Daten diskutiert.

II. FUNKENKAMMERN

1. Einleitung: Prinzip und Definition der wichtigsten Parameter

Heute ist die Funkenkammer eines der wichtigsten Nachweisgeräte der Elementarteilchenphysik. Obwohl bereits Ende der Vierzigerjahre die Bedingungen untersucht wurden, unter denen ionisierende Teilchen zwischen parallelen Elektroden Funken hervorrufen, dauerte es bis 1959, dass von S. Fukui und S. Miyamoto¹⁾ die erste, für Experimente brauchbare Funkenkammer gebaut wurde. 1960 wurde sie erstmals in Experimenten an Hochenergiebeschleunigern eingesetzt und ihre Brauchbarkeit erwiesen. Seither wurden, je nach den speziellen Erfordernissen der einzelnen Experimente, verschiedene Arten entwickelt. Ihnen allen gemeinsam ist die prinzipielle Anordnung, wie sie Abb. 1 zeigt:

Abb.1: Prinzip der Anwendung einer Funkenkammer.



Zähler 1 und 2 registrieren den Durchgang eines geladenen Teilchens und geben einen Spannungspuls an die Koinzidenzeinheit ab. Deren Ausgangssignal schaltet einen Pulsgenerator, der die Hochspannung an die Kammerelektroden legt. Der Wert dieser Spannung ist so eingestellt, dass die Ionen und Elektronen, die das Teilchen im Kammergas zwischen den Elektroden erzeugt hat, so beschleunigt werden, dass sie selbst ionisieren und soviel neue Ladungsträger erzeugen, dass sie entweder als Strom oder als sichtbarer Funke nachgewiesen werden können. Bereits mit Hilfe dieses einfachen Schemas lassen sich die wichtigsten Parameter einer Funkenkammer definieren:

Nachweiswahrscheinlichkeit (efficiency)

Die Nachweiswahrscheinlichkeit einer Kammer ist das Verhältnis der Anzahl der von der Kammer registrierten Teilchen zur Gesamtzahl der Teilchen, die die Kammer passiert haben. Besonders wichtig für die Anwendung bei Experimenten, bei denen Ereignisse mit vielen geladenen Teilchen registriert werden sollen, ist die Nachweiswahrscheinlichkeit für viele Teilchen gleichzeitig (multi-particle efficiency). Sie ist

definiert als das Zahlenverhältnis der Ereignisse, bei denen alle Teilchen registriert werden, zur Gesamtzahl der Ereignisse.

Gedächtniszeit (memory-time)

Verstreicht zwischen dem Durchgang eines Teilchens durch die Funkenkammer und dem Anlegen der Hochspannung an die Elektroden ein Zeitintervall, so geht ein Teil der erzeugten Ladungsträger durch Rekombination mit den Gasmolekülen oder Diffusion zu den Elektroden verloren, so dass die Nachweiswahrscheinlichkeit mit wachsender Verzögerungszeit t sinkt. Die Gedächtniszeit ist als jene Zeitspanne definiert, bei der die Nachweiswahrscheinlichkeit auf $\sim 95\%$ des Wertes für $t = 0$ gefallen ist. Sie kann durch Anlegen eines Saugfeldes (clearing-field), welches das Wandern der Ladungsträger zu den Elektroden beschleunigt, oder durch Zusätze elektronegativer Gase zum Kammergas (sie verkürzen die Lebensdauer der freien Elektronen) verkürzt werden. Sollen die Kammern in Teilchenstrahlen mit hohen Intensitäten arbeiten, so muss die Gedächtniszeit so kurz wie möglich sein, damit nicht zu viele Teilchen registriert werden, die zufällig vor oder nach einem Ereignis die Kammer durchquert haben. Sollen die Kammern eine gute Nachweiswahrscheinlichkeit besitzen, so ist ihre untere Grenze gegeben durch die

Ansprechzeit (resolution time),

die minimale Zeit, welche notwendig ist, um nach dem Durchgang eines Teilchens durch die Kammer die Hochspannung an die Elektroden anzulegen. Sie wird durch die Verzögerungszeit der Photovervielfacher, Kabel, Koinzidenzeinheiten und Schaltzeit des Hochspannungsgenerators bestimmt, und liegt in der Größenordnung 200-500 nsec.

Totzeit (dead-time)

Sie ist als jene Zeitspanne definiert, die mindestens verstreichen muss, bis die Kammer wieder zufriedenstellend arbeiten kann. Gegeben ist sie entweder durch die Ladezeit des Hochspannungspulsgenerators, oder durch die Lebensdauer freier Ladungsträger im Kammergas. Die Ladungsträgerdichte im Kammervolumen muss nämlich nach der Funkenbildung so weit absinken, dass sich beim Anlegen der Hochspannung nur dann ein Funke ausbilden kann, wenn vorher ein geladenes Teilchen die Kammer

durchquert hat.

Die Unterschiede zwischen den einzelnen Funkenkammertypen lassen sich am besten aus der Theorie der Gasentladung ableiten, auf die im Folgenden kurz eingegangen werden soll.

2. Theorie der Gasentladung^{2,3)}

Da in den meisten Funkenkammern hauptsächlich Edelgase als Füllgase verwendet werden, sollen hier nur die Vorgänge in Edelgasen und der Einfluss von kleinen Zusätzen elektronegativer Gase behandelt werden.

Ein minimal ionisierendes, geladenes Teilchen erzeugt beim Durchgang durch Neon von 1 atm. etwa 30 Ladungsträgerpaare pro cm. In der Zeit bis zum Anlegen der Spannung an die Elektroden (200-500 nsec) reduziert sich die Zahl der Ladungsträger durch folgende Prozesse:

- Rekombination:

In reinen Edelgasen ist die Rekombinationswahrscheinlichkeit vernachlässigbar klein. Geringe Beimengungen elektronegativer Gase verringern jedoch - wie die Gedächtniszeitmessungen als Funktion der Freonbeimengungen an unseren Kammern (Abschnitt IV/3) zeigen - sehr stark die mittlere Lebensdauer der freien Elektronen im Gas.

- Diffusion an die Elektroden.

Für die im Vergleich zu den Ionen bedeutend beweglicheren Elektronen ist der Diffusionskoeffizient D in Neon bei 1 atm. und 20°C etwa $10^3 \text{ cm}^2/\text{sec}$. Demnach beträgt nach etwa 500 nsec die mittlere quadratische Verschiebung der Elektronen $\overline{x^2} = 2 \cdot D \cdot t$ etwa 0.1 mm^2 . Deshalb ist der Prozentsatz der Ladungsträger, die bei einem Plattenabstand von mehr als 5 mm durch Diffusion zu den Elektroden verloren gehen, vernachlässigbar klein zu der Gesamtzahl der erzeugten Ladungsträger. Durch Anlegen eines Saugfeldes kann das Wandern der Ladungsträger zu den Elektroden beschleunigt werden.

Nach etwa 200-500 nsec wird ein Hochspannungspuls an die Elektroden gelegt. Wegen ihrer höheren Masse sind die Ionen bedeutend träger als die Elektronen und können in den folgenden Überlegungen als ruhend angesehen werden. Im angelegten Feld werden die Elektronen beschleunigt.

Jedes Elektron erzeugt nach dem Townsend'schen Gesetz α neue Elektronen pro cm. α ist der erste Townsend-Koeffizient; er ist eine Eigenschaft des Gases und eine Funktion des Quotienten E/p (E ist das elektrische Feld, das auf die Elektronen wirkt, p der Gasdruck).

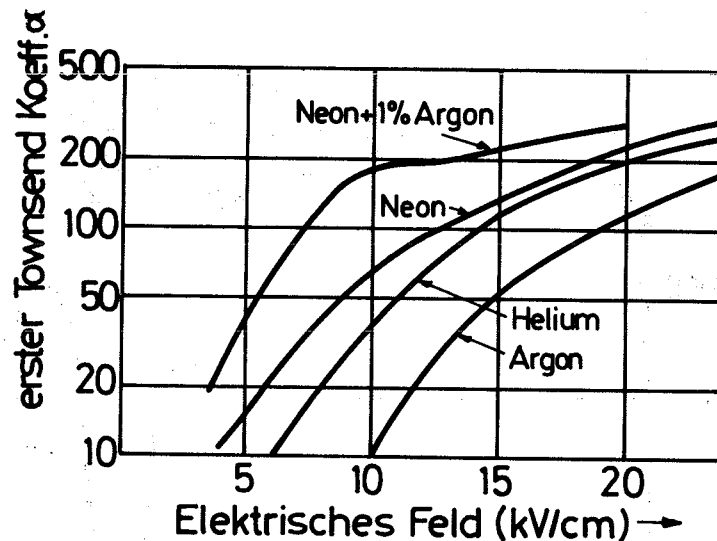


Abb. 2: Werte des ersten Townsend-Koeffizienten für Edelgase bei 1 atm.

Abb. 2 zeigt α für verschiedene Gase bei 1 atm. α ist für Neon und Argon jeweils kleiner, als für eine Mischung von Neon mit wenigen Prozenten Argon. Der Grund dafür ist, dass metastabile Neonatome, mit einem Anregungspotential von 16.5 V, durch Stoss Argonatome (Ionisierungspotential 15.68 V) ionisieren können. Experimentell wurde festgestellt, dass bei Atmosphärendruck eine Mischung von 1% Argon in Neon die niedrigste Durchbruchspannung für Ne-Ar Gemische ergibt.

Jedes Primärelektron bildet also eine Lawine von Sekundärelektronen, deren Kopf sich in Feldern, die notwendig sind, um einen Durchbruch zu erreichen, mit einer Geschwindigkeit von etwa 2×10^7 cm/sec auf die Anode zubewegt und einen Schwanz positiver Ionen hinter sich lässt. Je mehr neue Ladungsträger gebildet werden, desto mehr wird das ursprünglich homogene Feld gestört.

Abb.3: SCHEMATISCHER FELDVERLAUF EINER
ELEKTRONENLAWINE IM HOMOGENEN FELD
(E_0) EINER FUNKENKAMMER

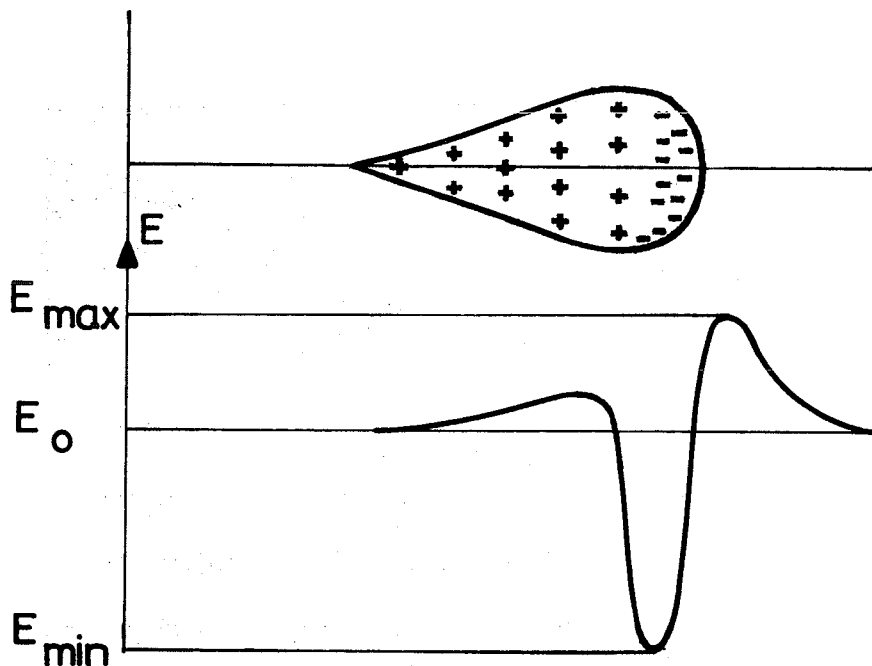


Abb. 3 zeigt schematisch Ladungsverteilung und Feldverlauf, während der Lawinenbildung. Mit zunehmender Elektronendichte im Lawinenkopf wird der Einfluss des äusseren Feldes innerhalb der Elektronenlawine immer mehr geschwächt und das Wachsen der Lawine gebremst, bis es, wenn sich beide Felder etwa kompensieren, zum Stillstand kommt. Falls das angelegte Feld so stark ist, dass sich bis zu diesem Augenblick im Kopf der Lawine etwa 10^8 Ladungsträger/mm³ gebildet haben, wird ein anderer Prozess dominierend: Die hohe Ladungsträgerdichte und die fallende Relativgeschwindigkeit zwischen Elektronen und positiven Ionen lassen die Wahrscheinlichkeit für Elektron-Ion-Rekombinationen unter isotroper Aussendung von Lichtquanten stark ansteigen. Die Photonen, deren mittlere freie Weglänge in Helium etwa 3-5 cm ist, setzen im Gebiet um die Hauptlawine Elektronen frei (Photo-

effekt). Im Bereich um die Enden der Lawine ist das Feld am stärksten, damit α am grössten und die Photoelektronen erzeugen dort neue Elektronenlawinen in Richtung des lokalen Feldes. Dieser Vorgang wird als Aufbau eines "streamers" (Kanalaufbau) bezeichnet. Bleibt das äussere Feld noch weiterhin angelegt, dann wächst der "streamer" durch den Photoeffekt, bis ein Plasmakanal die beiden Elektroden verbindet und sie als Funke entlädt. Abb. 4 (Seite 8) zeigt schematisch den Aufbau einer Funkenentladung in seinen einzelnen Stadien. Abb. 5 zeigt Photos von den einzelnen Stadien für den Fall, dass das ionisierende Teilchen senkrecht zum angelegten Feld die Kammer durchfliegt. Den zeitlichen Verlauf dieses Prozesses hat H. Tholl⁴⁾ photographisch in einem Stickstoff-Methangemisch bei 3 cm Elektrodenabstand untersucht und gefunden, dass die Lawinenbildung etwa 10mal so lange dauert wie die Zeit, in der sich aus der Lawine ein Streamer und der anschliessende Funke entwickelt.

3. Die verschiedenen Funkenkammern^{5,6)}

Je nachdem, zu welchem Zeitpunkt das äussere Feld abgeschaltet wird, kann der Prozess der Funkenbildung unterbrochen werden oder nicht; demnach unterscheidet man:

Streamer-Kammern

Spurfolgende ("track-following") Kammern

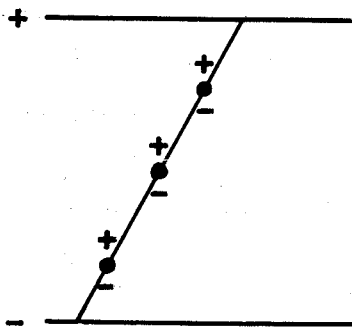
Die Eigenschaften beider sollen im Folgenden kurz beschrieben werden.

a) Streamer-Kammern

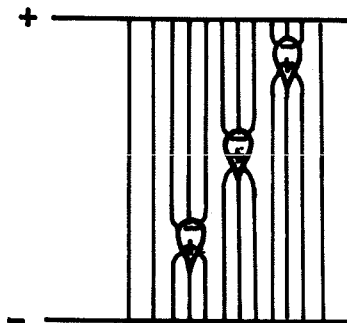
Durch eine Funkenstrecke parallel zur Kammer wird die Streamerbildung in jenem Augenblick unterbrochen, wenn die Entladung gerade hell genug ist, um photographiert zu werden. Das Bild von Abb. 5 (S. 9) ist ein solches Photo. Um jedes Primärelektron bildet sich ein kurzer Plasmaschlauch parallel zur Feldrichtung in der Kammer. Unabhängig von ihrer Richtung erscheint eine Teilchenbahn als eine Reihe von kleinen leuchtenden Punkten (~ 1 mm) bzw. Linien (3-5 mm) je nachdem, ob sie parallel oder senkrecht zum elektrischen Feld beobachtet wird. Eine Kammer dieser Art wurde erstmals von G.E. Chikovani und Mitarbeitern⁷⁾ 1964 gebaut. Sie kann, wie die Blaskammer, komplizierte Ereignisse mit vielen geladenen Teilchen registrieren. Bei ihr wird die lokal im elektrischen Feld gespeicherte Energie in Licht umgewandelt, so dass sich die verschiedenen Entladungen längs der einzelnen Teilchenspuren nicht beeinflussen. Dies unterscheidet sie von den anderen Funkenkammertypen,

Abb. 4: AUFBAU EINER FUNKENENTLADUNG

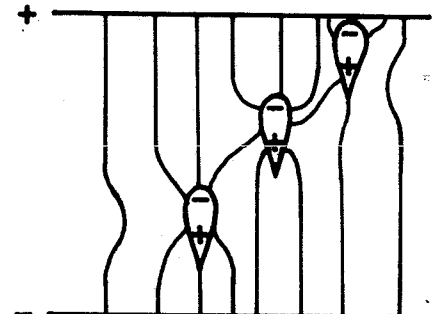
Aufbau der primären Elektronenlawinen



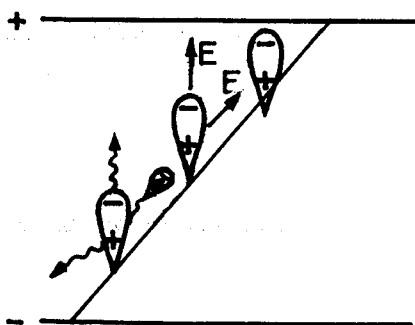
a) Ionisierung des Gases durch ein geladenes Teilchen



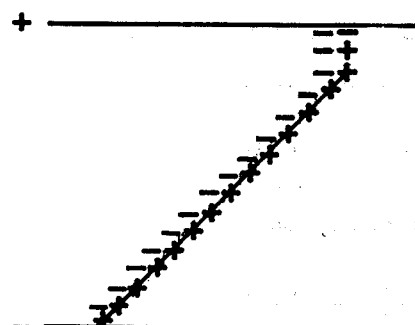
b) schwache Deformation des äußeren Feldes



c) starke



d) Aussendung von γ und Bildung neuer Lawinen im Bereich größter Feldstärke



e) Ausbildung des Funkenkanales und Durchschlag

Text zu Abb. 5:

Zeitliche Entwicklung eines Funkens in einer wide-gap Kammer, wenn die Teilchenbahn senkrecht zur Feldrichtung ist.

- Zeitpunkt zu dem sich aus den Elektronenlawinen Streamer bilden (zu diesem Zeitpunkt werden bei der Streamer Kammer die Elektroden kurzgeschlossen)
- Anwachsen der Streamer
- Anwachsen der Streamer
- Funkendurchbruch

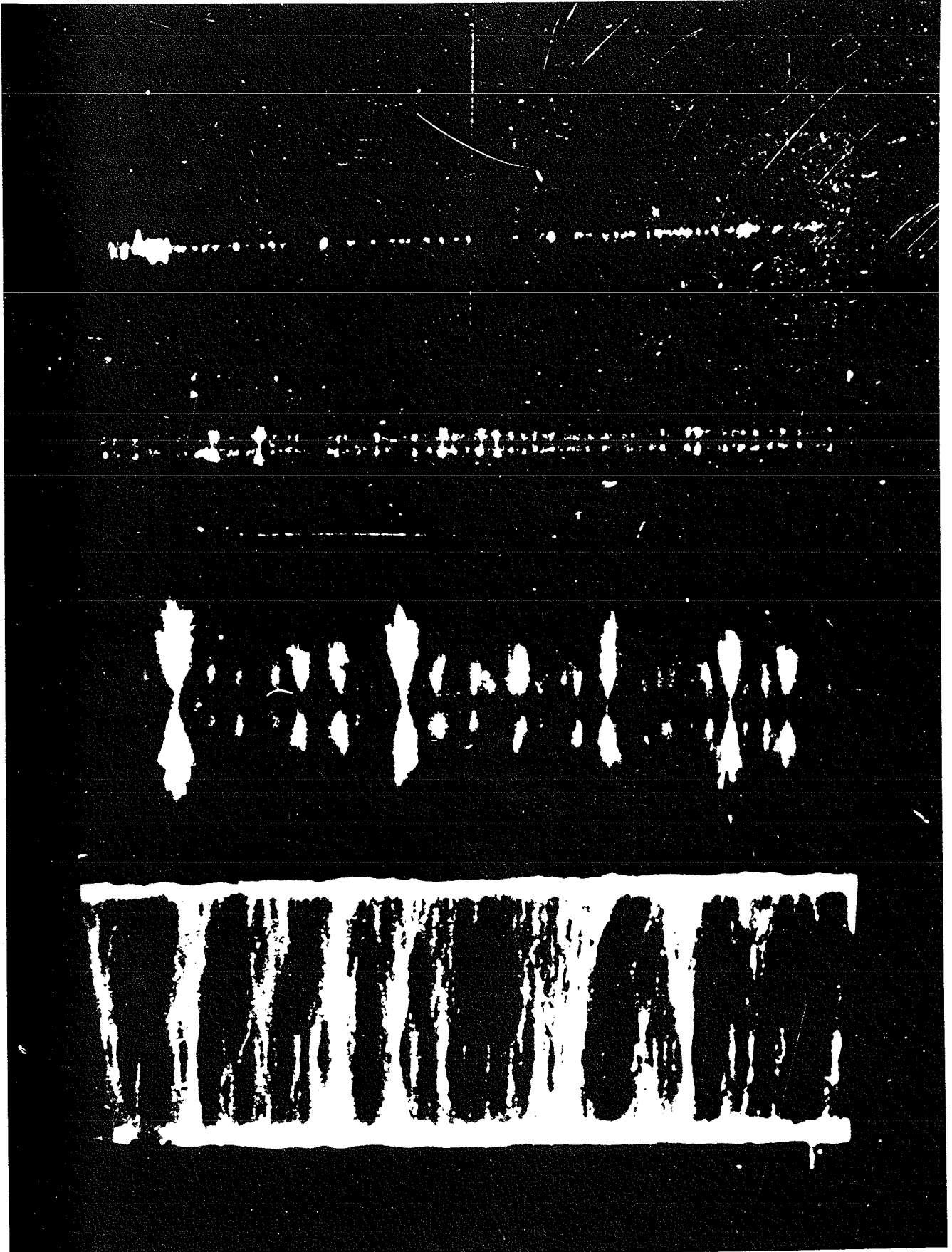


Abb. 5
(Text siehe Seite 8)

bei denen der durch die Plasmakanäle fliessende Strom optisch, akustisch oder elektrisch registriert wird. Ihr Vorteil gegenüber der Blaskammer ist ihre kürzere Totzeit (10 msec), sowie die Tatsache, dass sie mit Hilfe eines Systems von Zählern dann angeschaltet werden kann, wenn ein bestimmtes Ereignis vorliegt. Das räumliche Auflösungsvermögen von Streamerkammern ist zur Zeit um etwa einen Faktor 3 schlechter als das der Blaskammer. Bei kleinen Kammern ist es bereits gelungen, Teilchen mit verschiedenen Ionisationsdichten zu unterscheiden. Das Hauptproblem beim Bau der Streamerkammer ist weniger die Kammer selbst, als die Spannungsversorgung.

b) Kammern, bei denen die Funken den Teilchenbahnen folgen

Bei diesen Kammern wird die an den Elektroden angelegte Spannung nicht im Augenblick der Streamerbildung kurzgeschlossen, sondern die Kammer (Kapazität ~ 400 pF) im Wesentlichen durch den Plasmakanal entladen (bei vielen Kammern begrenzt ein Parallelwiderstand oder eine Parallelfunkenstrecke, welche nach dem Beginn der Entladung einschaltet wird, den Funkenstrom). Im Gegensatz zur "Streamerkammer" arbeiten sie nur, wenn die Teilchenspur ungefähr parallel ($< 45^\circ$, in Sonderfällen bis zu 60°) zum elektrischen Feld in der Kammer verläuft. Wie zuerst von A. Roberts et al.⁸⁾ gezeigt wurde, sind die elektrischen Eigenschaften einer spurfolgenden Funkenkammer weitgehend von der Konstruktion der Elektroden unabhängig. Die anschliessende Diskussion gilt unabhängig davon, ob die Elektroden aus Metallplatten, Metallfolien oder aus Drahtebenen gebildet sind. Drahtfunkenkammern haben sich für gewisse Methoden automatischen Auslesens sehr bewährt (siehe III.1b)).

Die spurfolgenden Kammern werden eingeteilt in "small"- und "wide-gap" Kammern. Obwohl sich "small-gap" und "wide-gap" Kammern zunächst nur durch den Elektrodenabstand ($\lesssim 3$ cm) unterscheiden, ist diese Einteilung dennoch auf Grund ihrer recht verschiedenen Eigenschaften gerechtfertigt.

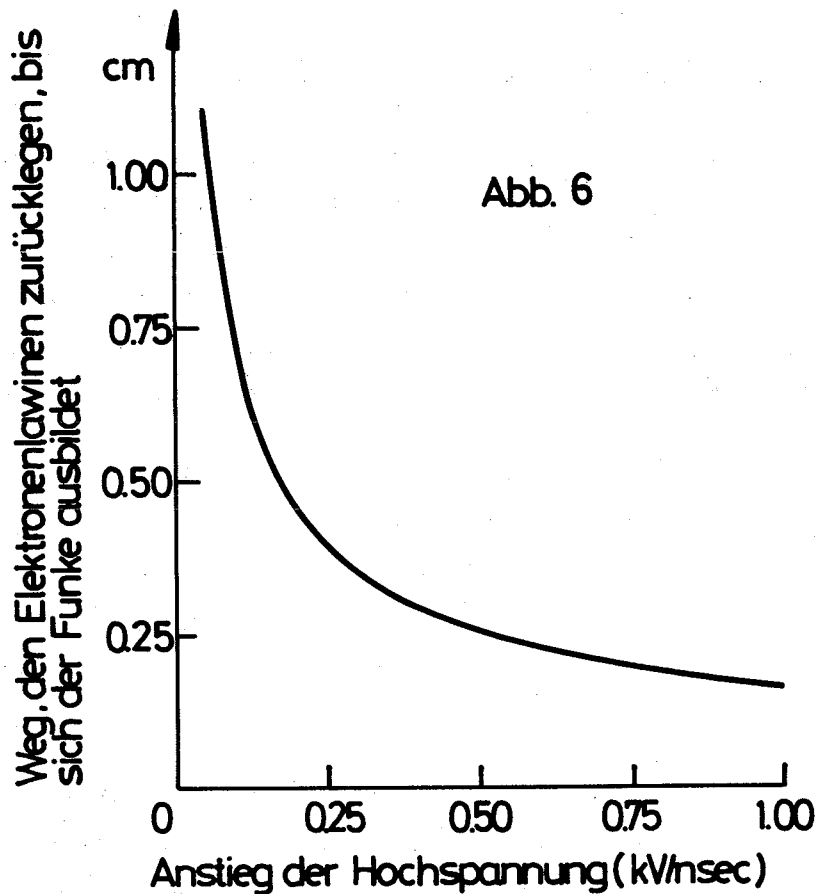
α) Funkenkammern mit kleinem Elektrodenabstand (small-gap)

In diesen Funkenkammern ist die Länge einer Elektronenlawine unmittelbar vor dem Einsetzen der Streamerbildung (für He etwa 0.25 cm bei einer üblichen Anstiegszeit der Spannung von 0.5 kV/nsec) vergleichbar mit dem Plattenabstand (8-10 mm), so dass der entstehende Funkenkanal nur aus wenigen Lawinen hervorgeht und damit, wie diese selbst, der Feldrichtung folgt. Daraus ergibt sich, speziell für gegen die Feldrichtung geneigte Bahnen, eine Abweichung des Funkens von der Teilchenspur. Eine genaue Bestimmung der Teilchenbahn wird durch Hintereinanderstellen von mehreren "small-gap" Kammern erreicht.

Die Gedächtniszeit lässt sich bei diesen Kammern mit Hilfe eines Saugfeldes leicht zwischen 20 μ sec (reines Edelgas, 0V/cm) und 0.5 μ sec (reines Edelgas, 500 V/cm) regulieren. Ein weiterer Vorteil ist die im Vergleich zu "wide-gap"-Kammern geringe Spannung von 5-10 kV. Sie lässt sich mit handelsüblichen Thyratrons mit Anstiegszeiten von etwa 20 nsec schalten.

Ein Problem bei diesen Kammern bildet die Nachweiswahrscheinlichkeit. Zunächst darf der Elektrodenabstand, damit man eine gute Nachweiswahrscheinlichkeit erhält, nicht zu klein sein. Wie in Abschnitt II,2 beschrieben wurde, müssen mindestens 10^8 Elektronen/mm³ in den Primärelektronenlawinen gebildet werden, damit aus ihnen ein Streamer und ein Funke entstehen kann. Der Weg, den die Elektronenlawinen bis zur Funkenbildung zurücklegen, muss natürlich kleiner sein als der Elektrodenabstand, da sonst alle Elektronen von der Anode aufgesammelt werden, noch bevor sich der Funke ausbilden kann. Abb. 6 (S.12) zeigt diesen Weg als Funktion der Anstiegszeit des Hochspannungspulses bei einem Endfeld von 10 kV/cm. Damit ergibt sich bei den üblichen Anstiegszeiten der Thyratrons von 0.5 kV/nsec eine untere Grenze von 3 mm für den Elektrodenabstand.

Das Problem der Nachweiswahrscheinlichkeit für mehrere Teilchen ist eng mit der Aufbauzeit t der Funken in der Kammer verbunden. Wird nämlich die Kammer bereits über einen Funken entladen, während sich ein anderer Funke noch ausbildet, so wird dessen vollständige Entwicklung wegen des Spannungsabfalles an den Kammerelektroden verhindert, und er kann nicht registriert werden. Die Entladung selbst



geschieht sehr schnell, da der Widerstand eines Funkens wegen des kleinen Elektrodenabstandes klein ist. Es gibt viele Ursachen, warum Funken sich verschieden schnell ausbilden können. Variiert der Elektrodenabstand d um den Betrag Δd , dann verursacht dies eine Feldänderung $\frac{\Delta E}{E} = \frac{\Delta d}{d}$. Für die Feldabhängigkeit von α in Neon gilt angenähert $\alpha \sim E^{7/2}$, so dass $\frac{\Delta \alpha}{\alpha} = \frac{7}{2} \frac{\Delta E}{E}$. Nach II,2 beschreibt das Townsend'sche Gesetz die Vorgänge im längsten Zeitintervall während der Funkenbildung. Um in beiden Feldern E und $E + \Delta E$ die gleiche Zahl N von Ladungsträgern zu erzeugen, gilt wenn die mittlere Driftgeschwindigkeit v unabhängig von E angesehen wird, was für kleine ΔE recht gut erfüllt ist:

$$N \sim e^{\alpha x} = e^{\alpha t v} = e^{(\alpha + \Delta\alpha)(t + \Delta t)v},$$

hieraus folgt für die Zeitdifferenz

$$\frac{\Delta t}{t} = - \frac{\Delta\alpha}{\alpha} = - \frac{7}{2} \frac{\Delta E}{E} = \frac{7}{2} \frac{\Delta d}{d}$$

Ein $\frac{\Delta d}{d}$ von 1% (d.h. bei einem Elektrodenabstand von 0.5 cm 5/1000 cm) bringt bereits einen Zeitunterschied von 3.5%, während sich der Übergang von Elektronenlawine zum Funkenkanal in 10% der Gesamtaufbauzeit vollzieht (Siehe II,2).

Statistische Schwankungen in der Anzahl der von dem durch die Kammer fliegenden Teilchen erzeugten Ladungsträger bringen ebenfalls verschiedene Aufbauzeiten mit sich. Dieses Problem wurde 1963 von Schneider untersucht⁹⁾. Seine Ergebnisse lassen sich mit der Formel $\frac{\Delta t}{t} = (20n^{1/2})^{-1}$ annähern, wobei n die Zahl der primär erzeugten Ladungsträger ist. Für eine Kammer mit 0.5 cm Elektrodenabstand ist n in Neon bei 1 atm. etwa 15 und damit $\frac{\Delta t}{t} \cong 1.3\%$.

Der Einfluss der Schwankungen der Funkenaufbauzeit lässt sich vermindern, wenn die Entladungszeit der Kammern verlängert wird.

β) Funkenkammern mit grossem Elektrodenabstand (wide-gap)

Bei diesen Kammern ist die Länge der Primärelektronenlawinen klein zur Funkenlänge, so dass die Funken bis zu grösserem Winkel den Teilchenbahnen folgen. Die Bedingungen, unter denen die grössten Winkel auftreten, lassen sich aus der Theorie der Funkenentladung ableiten (Seite 4). Einer Teilchenspurspur von grossem Neigungswinkel folgt der Funke umso besser, je grösser der Beitrag der Lawinen zum resultierenden elektrischen Gesamtfeld zwischen Kopf und Schwanz aufeinanderfolgender Elektronenlawinen ist. Dann werden die durch Photoionisation erzeugten Lawinen zwischen den Primärlawinen im Bereich der dort herrschenden grössten Feldstärke am schnellsten wachsen und so einen schrägen Funkenkanal bilden. Dies tritt dann bevorzugt ein, wenn die Lawinen kurz sind und ihre Ladungsträgerdichte gross ist. Dies wird durch eine kleine Anstiegszeit des Hochspannungspulses erreicht. Bei einer gegebenen Hochspannungsversorgung wird die Anstiegszeit durch Ladewiderstand und Kapazität der Kammer bestimmt. Wir erreichten bei einer Kammer mit 5 cm Elektrodenabstand und einer Fläche

von $40 \times 80 \text{ cm}^2$ Spuren bis zu 60° , während bei gleicher Spannungsversorgung Kammern mit 2 Elektrodenpaaren von $150 \times 150 \text{ cm}^2$ und 5 cm Elektrodenabstand nur 40° geneigte Spuren zeigten.

Ein weiterer Vorteil von Kammern mit grossem Elektrodenabstand ist ihre grosse Mehrteilchennachweiswahrscheinlichkeit. Bei ihnen ist es viel einfacher, grössere relative Schwankungen im Elektrodenabstand zu vermeiden. Auch ist die Zahl der erzeugten Ladungsträger grösser; damit sind die relativen statistischen Schwankungen kleiner, so dass die Schwankungen der Funkenbildungszeit abnehmen. Schliesslich ist wegen des grösseren Elektrodenabstandes der Funkenwiderstand grösser, die Entladungszeit damit länger, so dass Unterschiede zwischen den Aufbauzeiten von Funken weniger bedeutend werden. Wie gut der Vielteilchenwirkungsgrad sein kann, zeigen Messungen von Bolotov et al.¹⁰⁾ Die Autoren geben $\sim 100\%$ Nachweiswahrscheinlichkeit bei $\sim 10 \text{ cm}$ Elektrodenabstand für Ereignisse mit bis zu 50 Teilchen an. Ein Nachteil ist der grössere Aufwand bei der Erzeugung von Pulsen von einigen 10 kV mit kurzen Anstiegszeiten. Der Marxgenerator (siehe II,5) ist eine Lösung dieses Problems. Neuerdings gibt es auch Thyratrons, welche 80 kV in 25 nsec schalten.

Um die Gedächtniszeit der Kammern zu verringern, lässt sich wegen des grossen Elektrodenabstandes kein Saugfeld anwenden: Um bei einer Kammer mit 5 cm Elektrodenabstand eine Gedächtniszeit von etwa $1 \mu\text{sec}$ zu erreichen, wäre eine Elektronendriftgeschwindigkeit von $5 \text{ cm}/10^{-6} \text{ sec}$, und damit eine Feldstärke von etwa 5 kV/cm notwendig. Bei diesen Feldstärken ist α jedoch bereits so gross, dass im Gasvolumen mehr Ladungsträgerpaare gebildet, als von den Elektroden aufgesaugt werden. Dagegen lässt sich durch geringe Beimengungen elektronegativer Gase zum Kammergas die Lebensdauer freier Elektronen verkürzen und somit die Gedächtniszeit zwischen 0.5 und etwa $15 \mu\text{sec}$ regeln. Diese Methode hat sich bei unseren Kammern sehr bewährt.

4. Die gebauten Kammern¹¹⁾

Im beschriebenen Experiment wurden ausschliesslich Drahtfunkenkammern verwendet. Von ihnen handeln die nächsten Abschnitte.

Es wurden zwei verschiedene Typen von wide-gap-Kammern gebaut. (Abb. 7 zeigt einen Prototyp der gebauten Kammern mit etwas verschiedenen Abmessungen):

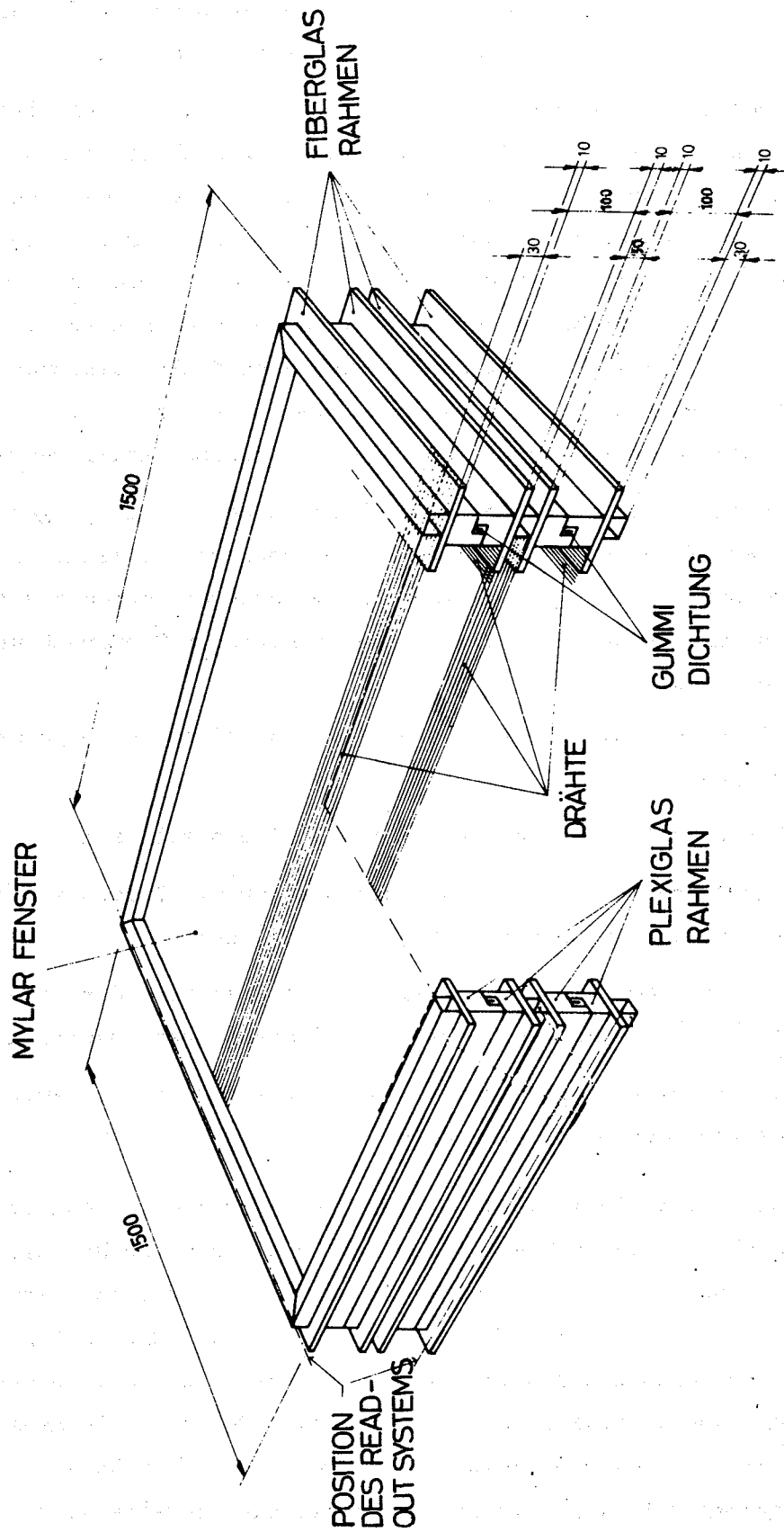


Abb. 7: SCHEMATISCHE DARSTELLUNG EINER GROSSEN KAMMER

- Eine Zweielektrodenkammer mit einer Fläche von 40 cm x 80 cm und 5 cm Elektrodenabstand (eine Ebene parallel gespannter 70 μ dicker Kupfer-Berylliumdrähte in 1 mm Abstand bildet eine Elektrode).
- Eine Vierelektrodenkammer von 150 cm x 150 cm und 5 cm Elektrodenabstand. Zur Verminderung der elektromagnetischen Abstrahlung beim Funkendurchbruch wird die Hochspannung an die beiden inneren Ebenen gelegt, während die äusseren geerdet sind. Um zwei Koordinaten pro Elektrodenpaar auslesen zu können, verlaufen die Drähte in den beiden Ebenen eines Elektrodenpaares senkrecht zueinander. Um bei Ereignissen mit mehreren Funken einander zugeordnete Koordinatenpaare eindeutig aufzufinden, sind die Elektrodenpaare als ganze um einen Winkel von 10° gegeneinander gedreht.

Die Rahmen beider Kammern wurden aus Fiberglas und Plexiglas gebaut; die Fenster aus Mylar von 50 μ Dicke - ein Kompromiss zwischen minimaler Vielfachstreuung für geladene Teilchen beim Durchgang durch die Kammer, ausreichender mechanischer Festigkeit gegen geringen Überdruck und Dichtheit gegenüber He-Gas.

Bei der Wahl von Drahtdurchmesser und Drahtabstand waren folgende Überlegungen massgebend:

- Die Vielfachstreuung in den Drähten soll minimal sein.
- Der Widerstand der Drähte soll klein gegen den Widerstand der Funkenstrecke beim Durchbruch ($\approx 20 \Omega/\text{cm}$) sein.
- Das elektrische Feld in der Kammer soll ausreichend homogen sein, so dass sich die Funken längs der Teilchenbahnen ausbilden können.
- Der Drahtabstand soll nicht grösser als 1.5 mm sein, damit das Auflösungsvermögen nicht schlechter als etwa ± 1 mm ist.

Wie experimentell bestimmt wurde, liegen bei 5 cm Elektrodenabstand die Grenzen für minimalen Drahtdurchmesser und maximalen Drahtabstand etwa bei 60 μ bzw. ~ 1.5 mm. Bei kleinerem Drahtdurchmesser ist das elektrische Feld, das im Bereich um die Drähte etwa umgekehrt proportional zum Abstand vom Drahtzentrum ist, so gross, dass sich beim Anlegen der Hochspannung Funken bilden, ohne dass geladene Teilchen die Kammer durchquert haben. Bei grösserem Abstand zwischen den Drähten ist die Zone inhomogenen Feldes im Bereich der Elektroden bereits so gross,

dass die Ausbildung von Funken längs der Teilchenbahnen gestört wird und nur ein sehr schlechtes Auflösungsvermögen erreicht werden kann.

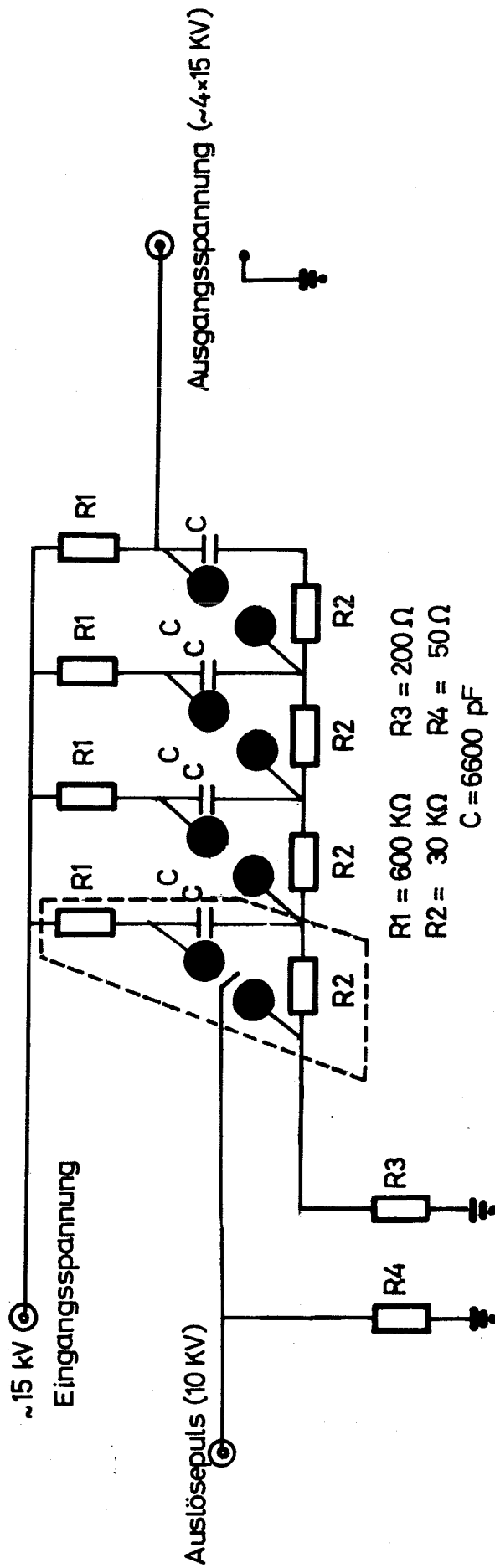
5. Die Spannungsversorgung^{1,2)}

Der Hochspannungspuls wird von einem vierstufigen Marx-generator erzeugt. Abb. 8 (Seite 18) zeigt den Schaltplan des verwendeten Generators. Eine Stufe ist in Abb. 8 durch die schraffierte Umrandung angedeutet. Die Elektroden der Funkenstrecken bestehen aus Stahlkugeln, ihr Abstand ist regulierbar. Die Funkenstrecken sind so angeordnet, dass UV-Licht, das beim Durchbruch einer Funkenstrecke entsteht das Gas (in unserem Falle Luft) in den anderen Funkenstrecken ebenfalls ionisiert.

Im Folgenden soll die Arbeitsweise des Generators erklärt werden: Zunächst werden die Kapazitäten C über die Widerstände R1 parallel aufgeladen (Ladezeit etwa 6 msec). Damit liegt an jeder Funkenstrecke die Eingangsspannung, deren Wert kleiner als die Durchbruchsspannung der Funkenstrecke sein muss. Die Entladung des Generators wird durch den 10 kV-Puls eines Thyratrons ausgelöst. Dieser erzeugt einen Funken in der ersten Funkenstrecke und bringt diese zum Durchbruch. Die dabei entstehenden Lichtquanten ionisieren das Gas in anderen Funkenstrecken und bringen diese ebenfalls zum Durchbruch. Der Widerstand der Funkenstrecken ist nun vernachlässigbar klein gegenüber R1 und R2. Dadurch sind nun die Kapazitäten C nicht mehr parallel, sondern in Reihe geschaltet und die Ausgangsspannung ist (bei grossem Verbraucherwiderstand) gleich dem Produkt Eingangsspannung und Zahl der Stufen. Nun entladen sich die Kapazitäten über den Verbraucher bis die Spannung an den Funkenstrecken unter die Löschspannung gefallen ist und die Kapazitäten erneut über die Widerstände R1 parallel aufgeladen werden.

Mit dem gebauten Generator erhält man - wenn die Funkenstrecken auf die gleiche Durchbruchspannung adjustiert sind - beim Aufladen einer grossen Kammer mit 4 Drahtebenen (900 pF) über einen 10 Ω Ladewiderstand, Anstiegszeiten von 20 nsec. Das Zeitintervall zwischen der Ankunft des Auslösesignals vom Thyatron und dem Beginn des Hochspannungssignals (15% der Endamplitude) beträgt 100 nsec mit einer Unsicherheit von 10 nsec. Die Totzeit des Generators ist kleiner als 8 msec.

Abb. 8: 4-Stufiger Marx-Generator



III. AUSLESESYSTEME VON FUNKENKAMMERN

1) Die verschiedenen Systeme^{6, 7)}

a) Photographische Methoden

Bei den ersten Funkenkammerexperimenten wurden die Funken ausschliesslich photographisch registriert. Die Helligkeit der Funken kann durch Parallelschaltung einer Kapazität zu den Elektroden gesteuert werden. Bei Ereignissen mit mehreren Spuren kann, besonders bei verschieden geneigten Teilchenbahnen, die Helligkeit der einzelnen Funken sehr verschieden sein, so dass entweder nicht alle Teilchen gelesen werden oder die hellen Funken wegen Überbelichtung nur ungenau gemessen werden können. Die Messgenauigkeit dieser Methode wird mit ± 0.2 mm angegeben.

Solange Funkenkammern verwendet werden, um seltene Ereignisse zu identifizieren, ist der Aufwand, die genommenen Bilder zu "scannen" und zu messen, nicht zu gross. Deshalb kann diese Technik nur in solchen Experimenten angewandt werden. Sollen jedoch Experimente mit Millionen von Ereignissen durchgeführt werden, dann wird der Aufwand der Auswertung zu gross, und man ist auf Methoden angewiesen, die die Information der Funkenkammern unmittelbar während das Experiment läuft digitalisiert speichern.

b) Digitalisierte Methoden

Diese Methoden lassen sich, je nachdem ob die Kammerelektroden aus Drahtebenen oder Metallfolien bestehen in zwei Gruppen einteilen.

a) Auslesesysteme für Kammern mit Plattenelektroden

In diese Gruppe fällt zunächst, als Weiterentwicklung der photographisch-optischen Methode, die Anwendung der direkt digitalisierenden Fernsehkamera¹³⁾. Bei dieser Methode wird das Bild des Funkens mit einer Fernsehkamera auf die empfindliche, halbleitende Schicht eines Vidicons fokussiert. Diese Schicht kann das Bild des Funkens als Potentialverteilung etwa 100 msec speichern.

Ein Elektronenstrahl, dessen Ablenkung in einem Kondensator sehr genau durch einen Oszillator gesteuert wird, fragt die empfindliche Videonschicht ab. Jedesmal wenn er auf Potentialunterschiede trifft, erzeugt er ein Videosignal, das auf Magnetband geschrieben wird und zusammen

mit der Phase des Ablenkoszillators die Information über die Funkenposition enthält. Die Genauigkeit, mit der damit Funken registriert werden können, beträgt $1/2500$, d.h. bei einer 1 m Kammer etwa 0.4 mm.

Ein weiteres System, das bei Kammern, deren Elektroden aus Metallfolien bestehen angewendet wird, ist das

akustische System.¹⁴⁾

Beim akustischen System wird mit Hilfe von Mikrofonen in den Ecken der Kammer die Laufzeit der durch den Funken erzeugten Schallwelle zwischen Funke und Mikrophon gemessen. Die Messgenauigkeit beträgt etwa ± 0.4 mm. Allerdings verhindert die geringe Dämpfung der Mikrophonmembran, dass eine kurz nach der ersten eintreffende Schockwelle von dieser getrennt werden kann. Dies begrenzt die Mehrteilchennachweiswahrscheinlichkeit und das räumliche Auflösungsvermögen.

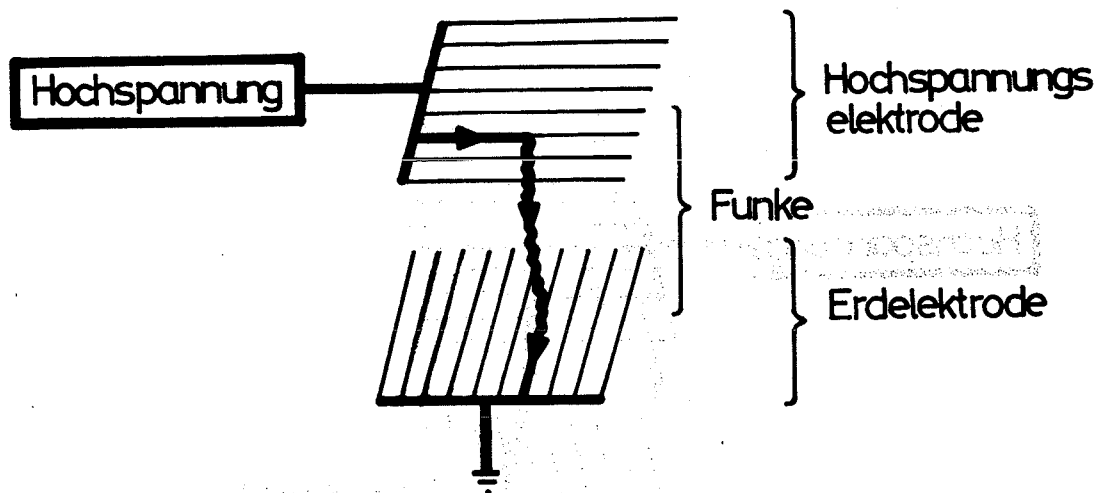
$\beta)$ Auslesesysteme für Kammern, deren Elektroden aus Drahtebenen gebildet sind

Bestehen die Elektroden einer Funkenkammer aus Drahtebenen, dann fließt der Strom beim Funkendurchbruch über einen Draht (oder mehrere benachbarte Drähte) der Hochspannungselektrode, den Funkenkanal und einen Draht der Erdelektrode von der Hochspannung zur Erde. Dies ist in Abb. 9 schematisch dargestellt. Aufgabe des Auslesesystemes ist es, jenen Draht zu finden, durch den der Strom fließt und damit die Koordinaten des Funkens festzustellen. Dies kann durch folgende Anordnungen erreicht werden:

Auslesen mit Ferritkernen¹⁵⁾

Bei dieser Methode werden die Elektrodendrähte einzeln durch Speicherkerne hindurchgeführt. Entlädt sich die Kammer, dann fließt durch jene Elektrodendrähte, auf die die Funken auftreffen, ein Strom, der die Magnetisierung der entsprechenden Speicherkerne ändert. Anschliessendes Abfragen der Speicherkerne erlaubt es, die Position der stromführenden Drähte festzustellen und damit die Koordinaten der Funken zu finden. Die Messgenauigkeit dieser Methode wird mit ± 0.3 mm bis ± 0.6 mm bei 1 mm Drahtabstand angegeben.

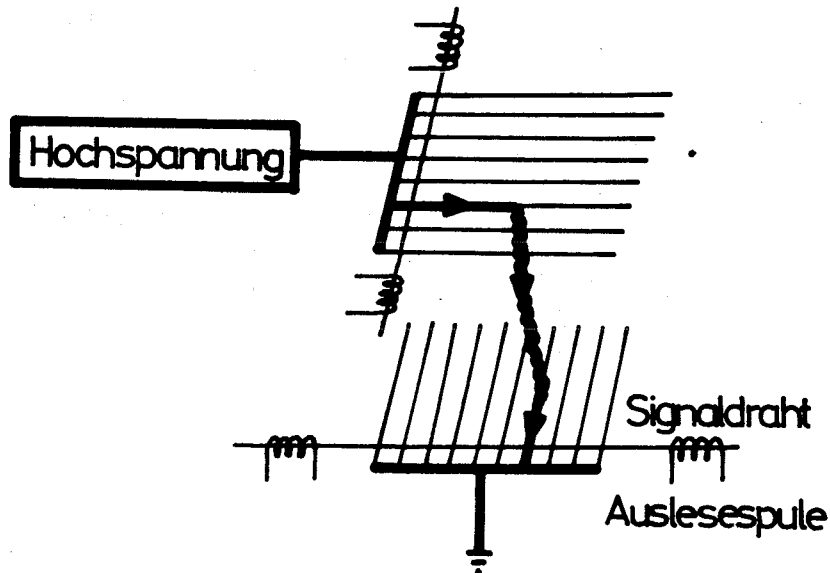
Abb.9: Schematische Darstellung einer Drahtfunkenkammer



Magnetostriktives System¹⁶⁾

Die Anordnung dieses Auslesesystems ist schematisch in Abb. 10 dargestellt: ein magnetostriktiver Signaldraht wird (elektrisch isoliert) senkrecht an den Drähten einer Kammerelektrode vorbeigeführt. Das Magnetfeld, das der Funkenstrom um einen Elektrodradht erzeugt, bewirkt im Signaldraht eine lokale Magnetostraktion und damit eine lokale Dichteänderung. Diese Dichteänderung pflanzt sich als akustische Welle längs des Drahtes fest. Wegen des inversen magnetostriktiven Effektes ist die Welle mit einem Magnetfeld gekoppelt, das in einer um das Ende des Signaldrahtes gewickelten Auslesespule einen Spannungspuls induziert. Die Koordinate des stromführenden Drahtes kann, bei bekannter akustischer Geschwindigkeit im Auslesedraht, aus der Laufzeit des magnetostriktiven Signales (Zeit zwischen Funkendurchbruch und Ankunft der Welle in der Auslesespule) bestimmt werden. Dieses System wurde bei den gebauten Kammern verwendet und wird in III,2 und III,3 noch genauer behandelt.

Abb.10: Schematische Darstellung einer Drahtfunkenkammer mit magnetostruktivem Auslesesystem



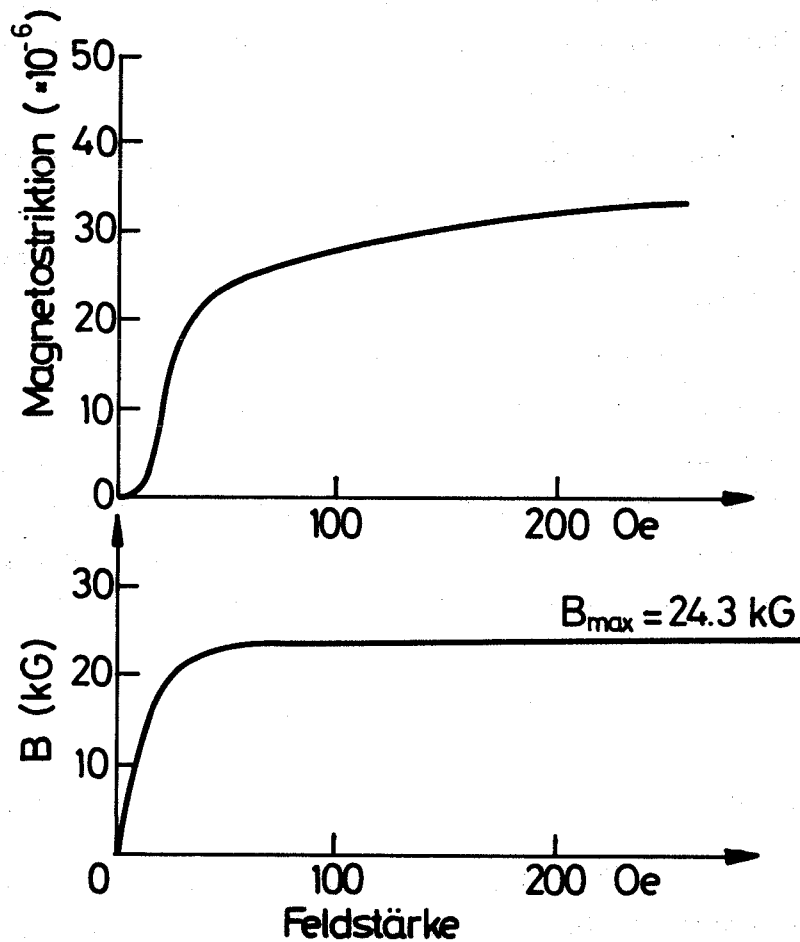
Sparkostruktives System¹⁷⁾

Dieses Auslesesystem ist dem magnetostruktiven sehr ähnlich: an Stelle des magnetostruktiven Auslesedrahtes wird jedoch ein gewöhnlicher Metalldraht, in geringem Abstand an den Elektrodenröhren senkrecht vorbeigeführt. Fließt der Funkenstrom durch einen Elektroden- draht, so schlägt ein Funke zu dem Signaldraht über und erzeugt in diesem eine mechanische Stosswelle. Mit Hilfe eines piezoelektrischen Kristalles, der am Drahtende befestigt ist, wird die mechanische Schwingung in einen elektrischen Puls umgewandelt. Aus der Laufzeit der Welle im Draht wird die Position des stromführenden Elektroden- drahtes ermittelt. Dieses System hat den Vorteil, auch in Magnetfeldern zu arbeiten.

2. Das magnetostruktive Signal

Abb. 11 (S. 23) zeigt die magnetische Induktion und darüber die relative magnetostruktive Längenänderung λ eines Co- Fe-Drahtes als Funktion des angelegten Magnetfeldes¹⁸⁾. Im linearen Teil der Hysteresekurve, wird der Draht dadurch magnetisiert, dass durch Wand-

Abb 11: Magnetostraktion und magnetische Induktion als Funktion der Feldstärke für Co/Fe 50/50



verschiebungen jene Weiss'schen Bezirke anwachsen, deren Magnetisierung parallel zum angelegten Feld ist. Dort ist λ klein. Im Bereich des Knies der Magnetisierungskurve wird die Probe hauptsächlich durch das Drehen der Weiss'schen Bezirke magnetisiert. Dort steigt λ stark an, bis es im Bereich der Sättigung einen konstanten maximalen Wert erreicht. Je nachdem, ob die Länge der Probe längs der Magnetisierung zu- oder abnimmt, unterscheidet man positive (Fe, Fe-Co, $\lambda > 0$) und negative (Ni, $\lambda < 0$) Magnetostriktion.

Der nun folgende Abschnitt soll die verschiedenen Vorgänge vom Entladen der Kammer durch den Funken bis zum Empfang des Spannungspulses an der Auslesespule beschreiben, um die Form der Signale und die

Bedingungen, unter denen die besten Ergebnisse erzielt werden, zu erklären (Bei Abschätzungen werden die an unseren Kammern gemessenen Werte eingesetzt).

Während sich die Kapazität des Hochspannungsgenerators und der Kammer über einen Funken entlädt, fliesst durch den vom Funken getroffenen Elektrodendraht etwa 250 nsec lang ein Strom von etwa 500 A.

Im folgenden wird als zeitliche Stromverteilung der Einfachheit halber ein Rechteckimpuls von etwa 500 A und 250 nsec angenommen. Die genaue Stromverteilung ist für die folgenden prinzipiellen Überlegungen nicht wesentlich.

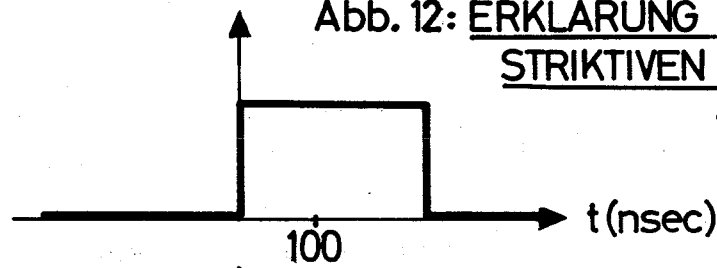
Der stromführende Elektrodendraht wird, im Abstand von etwa $\frac{1}{2}$ mm senkrecht an dem magnetostriktiven Draht elektrisch isoliert vorbeigeführt (Abb. 10). Wird der Strom (I) eingeschaltet, dann bildet sich längs des magnetostriktiven Drahtes (X ist die Koordinate längs des magnetostriktiven Drahtes, X = 0 ist der Kreuzungspunkt mit dem stromführenden Draht) ein magnetisches Feld:

$$H(X) = \frac{I}{2\pi(X^2 + d^2)^{1/2}}$$

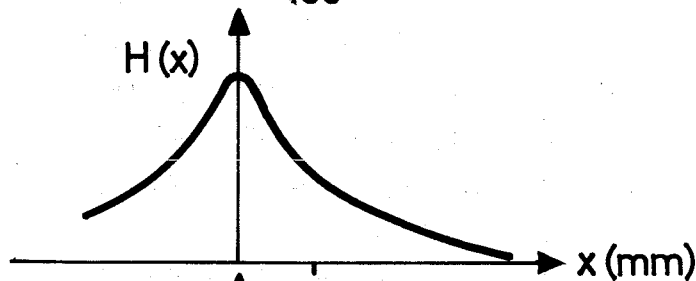
aus (d ist der Abstand Elektrodendraht - Auslesedraht).

Abb. 12b zeigt diese Feldverteilung für $d = \frac{1}{2}$ mm. Die Relaxationszeit (Relaxationszeit = Mass für die Zeit, die Weiss'schen Bezirke benötigen sich in einem äusseren Magnetfeld auszurichten) der Weiss'schen Bezirke beträgt weniger als 10 nsec¹⁸). In dieser Zeit kann sich eine Störung innerhalb des Drahtes nur etwa 5×10^{-2} mm fortpflanzen, was gegen die Breite der Feldverteilung (≈ 2 mm) klein ist. Damit ergibt sich für die durch die magnetostriktive Längenänderung bewirkte lineare Dichteänderung (g/mm) längs des Drahtes: $\Delta \rho(X) \sim \lambda(H(X))$. Mit λ aus Abb. 11 ergibt sich die Dichteverteilung von Abb. 12c. Etwa 250 nsec später, kurz vor Abschalten des Stromes, ist die Welle ungefähr 1.25 mm von der Stelle X = 0 weggewandert, so dass sich die Dichteverteilung von Abb. 12d ergibt. Schalten wir nun den Strom und damit das äussere Feld ab, so kehren die Weiss'schen Bezirke in einer Zeit, deren Grössenordnung durch die Relaxationszeit gegeben ist, in ihre Unordnung zurück.

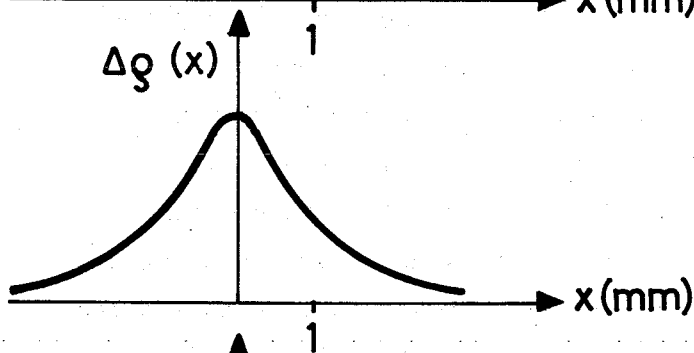
Abb. 12: ERKLÄRUNG DES MAGNETO-STRIKTIVEN SIGNALS



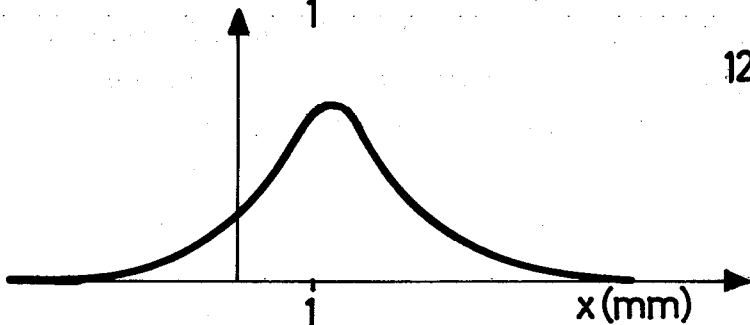
12 a) Strompuls



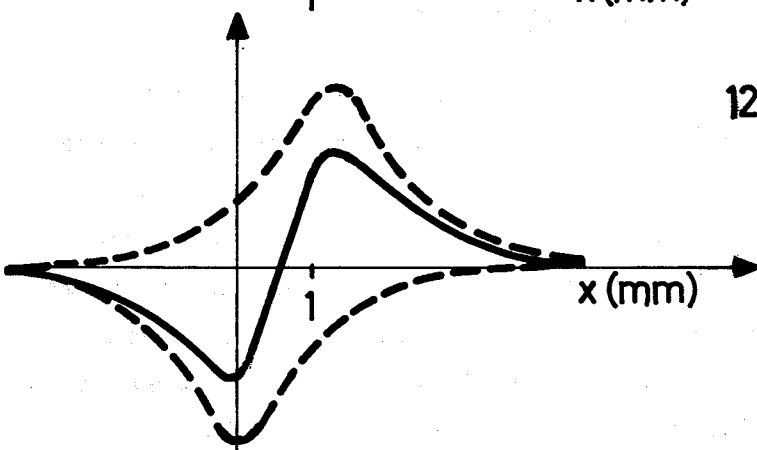
12 b) Magnetfeld an Drahtoberfläche unmittelbar nach Anschalten des Stromes



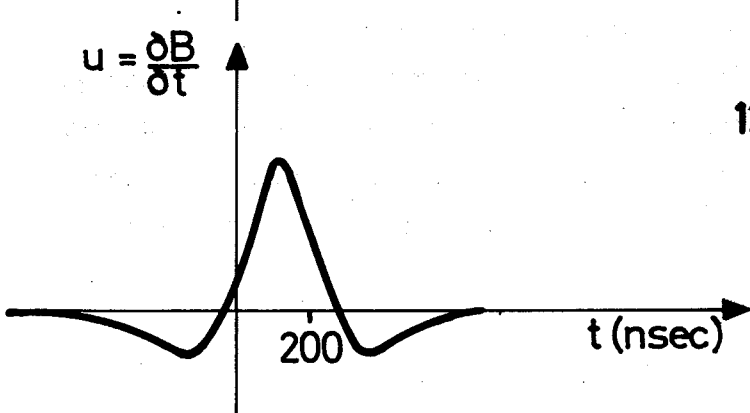
12 c) Dichteverteilung unmittelbar nach Anschalten des Stromes



12 d) Dichteverteilung unmittelbar vor Abschalten des Stromes (nur nach rechts laufende Welle gezeichnet)



12 e) Dichteverteilung unmittelbar nach Abschalten des Stromes



12 f) magnetostriktives Signal (Ableitung der letzten Kurve)

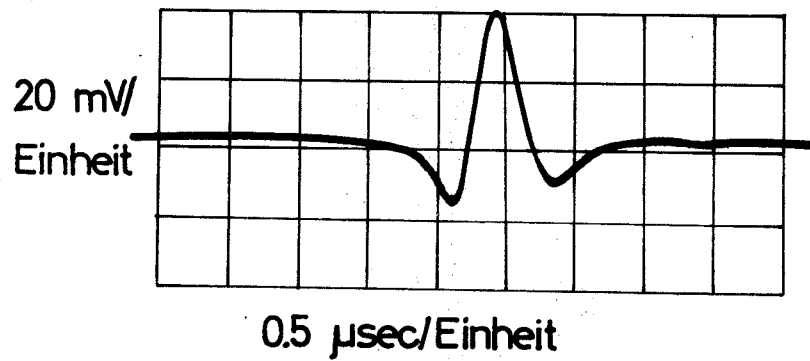
Dadurch zieht sich der Draht um $X=0$ lokal zusammen, und es bildet sich die S-förmige Dichteverteilung von Abb. 12e. (Sie ergibt sich als die Summe der beiden gestrichelten Kurven von Abb. 12e. Die obere gestrichelte Kurve ist die Kurve von Abb. 12d, also die lineare Dichteverteilung, die durch das Einschalten des Stromes zum Zeitpunkt $t=0$ entsteht; die untere gestrichelte Kurve ist die lineare Dichteverteilung, die durch das Abschalten des Stromes bei $t=250$ nsec entsteht. Sie ist die Kurve von Abb. 12c mit negativem Vorzeichen). Diese Dichteverteilung läuft als akustische Welle durch den Draht. Bei diesen magnetostriktiven Wellen handelt es sich im Wesentlichen um Oberflächenwellen, da eine Welle im Inneren des Drahtes durch Wirbelstromverluste gedämpft wird. Wegen der magnetostriktiven Eigenschaft des Drahtes sind Dichteunterschiede mit Magnetfeldern gekoppelt. Mit Hilfe einer um den Draht gewickelten Auslesespule, deren Länge kurz gegen die Signalbreite ist, kann daher die zeitliche (t) Änderung des mit der fortlaufenden Welle gekoppelten Magnetfeldes als Spannungspuls gemessen werden. Da $X=v \cdot t$ ist (v ist die Schallgeschwindigkeit im Draht) ergibt sich die Pulsform als Ableitung der Kurve 12e; sie ist in Abb. 12f gezeichnet. Abb. 13 zeigt zum Vergleich ein gemessenes Signal.

3. Das verwendete Auslesesystem

Anordnung des magnetostriktiven Drahtes (Abb. 14)

Als magnetostriktiver Draht wurde ein 50/50 Co-Fe-Vacofluxband (0.05×0.5 mm) verwendet. Massgebend für die Wahl war die grosse magnetostriktive Längenänderung der Co-Fe Legierung und die relativ kleine Dämpfung des Signales (Faktor 2 über 2m). Die Länge des Signaldrahtes beträgt etwa 2.20 m. Der Draht wird von einem Aluminiumsupport gehalten; mit Hilfe von Stellschrauben kann er gegen die Elektroden-drähte gedrückt werden. Da es sich bei den Wellen im Wesentlichen um Oberflächenwellen handelt, können sie mit Hilfe eines elastischen Silikonkautschuks an den Enden gedämpft werden (etwa ein Faktor 100 pro 10 cm). Somit werden störende Reflexionen vermieden. 20 Windungen eines lackierten Kupferdrahtes von 0.06 mm Durchmesser jeweils, direkt um den magnetostriktiven Draht gewickelt, dienen an beiden Enden des Drahtes als Auslesespulen. Jede Spule befindet sich im Magnetfeld

Abb.13: Magnetostruktives Signal
(Ausgangssignal des
Vorverstärkers)



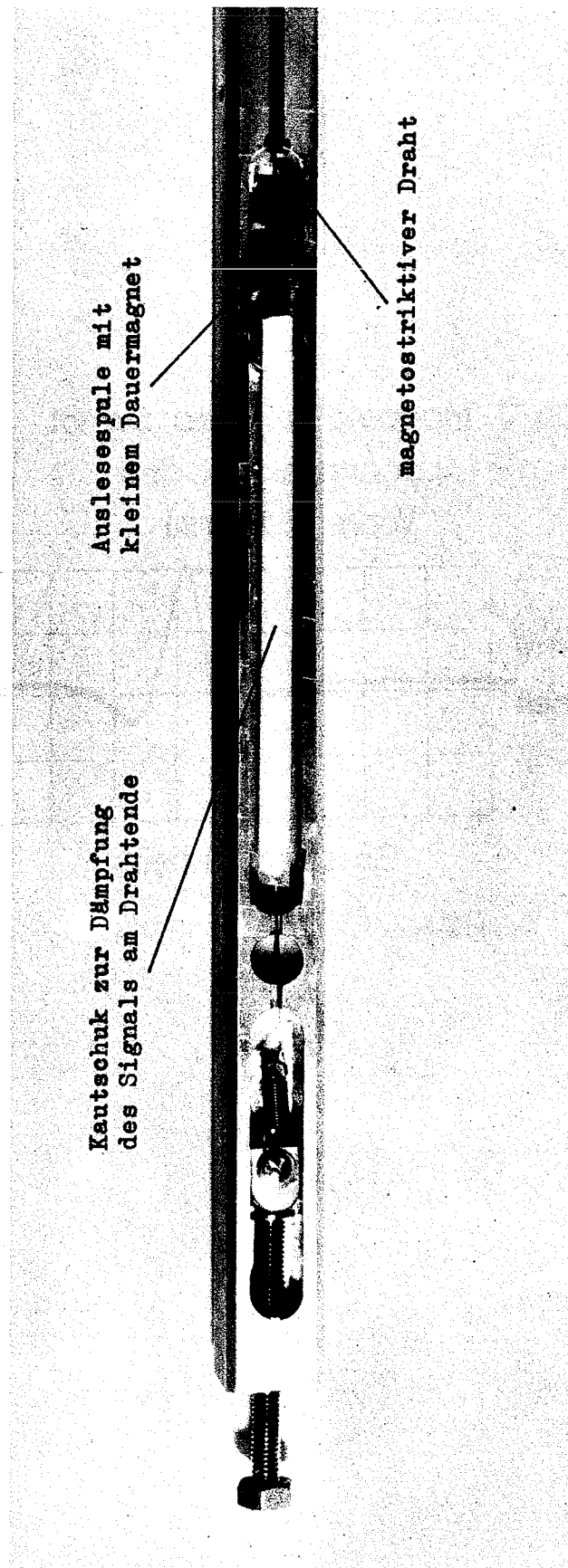


Abb. 14.
Magnetostruktive Auslesevorrichtung

eines kleinen Dauermagnetens, das den Draht nahezu senkrecht durchsetzt. Dadurch ist der Draht in Querrichtung magnetisiert, so dass eine durchlaufende magnetostriktive Welle maximale Flussänderung und damit maximale Amplitude des Ausgangssignals bewirkt. Die Richtung der kleinen Komponente längs des Drahtes bestimmt das Vorzeichen des Signals.

Als Schwierigkeit hat sich erwiesen, dass die Schallgeschwindigkeit im magnetostriktiven Draht sehr von dessen Oberflächenbeschaffenheit abhängt, die wegen der Herstellungsart (Ziehen) nicht immer ausreichend gleich war. Deshalb wurde für jeden Draht die Geschwindigkeit - in einigen Fällen variierte sie längs eines Drahtes bis zu 2% - gemessen, um genaues Auslesen zu erreichen.

Ausleseelektronik

Abb. 15 zeigt schematisch die elektronische Anordnung, mit der die Kammern ausgelöst und die Information über die Position des Funkens in der Kammer registriert werden. Die Funktion dieser Elektronik lässt sich am einfachsten am zeitlichen Ablauf eines Ereignisses erklären (als Experiment kann man sich dazu die Anordnung von Abb. 1, Seite 2 vorstellen). Die beiden Szintillationszähler registrieren den Durchgang eines geladenen Teilchens und geben einen Spannungspuls an die Koinzidenzeinheit ab. Treffen die beiden Pulse in der gewünschten Zeitfolge ein, dann definiert die Koinzidenzeinheit ein Ereignis und gibt ein Signal an die Kommandereinheit ab. Sie ist die zentrale Einheit, die alle folgenden Vorgänge steuert.

Unmittelbar bei Eintreffen des Eingangssignales (Zeitpunkt $t = 0$) gibt sie ein Signal an das Thyatron ab, welches durch einen 10 kV-Puls die Entladung des Marxgenerators auslöst (Kapitel II,5 Seite 17). Damit entlädt sich die Kammer über einen Funken längs der Teilchenspur, der im magnetostriktiven Draht eine akustische Welle erzeugt (Kapitel III,2).

Nach einem festen Zeitintervall Δt gibt der Kommander ein Signal an die Pulszähler ab, die mit einer 20 MHz-Uhr verbunden sind, so dass sie zum Zeitpunkt $t = \Delta t$ zu zählen beginnen. Ihr Zählerinhalt geteilt durch 20×10^6 zeigt somit die seit dem Zeitpunkt $t = \Delta t$

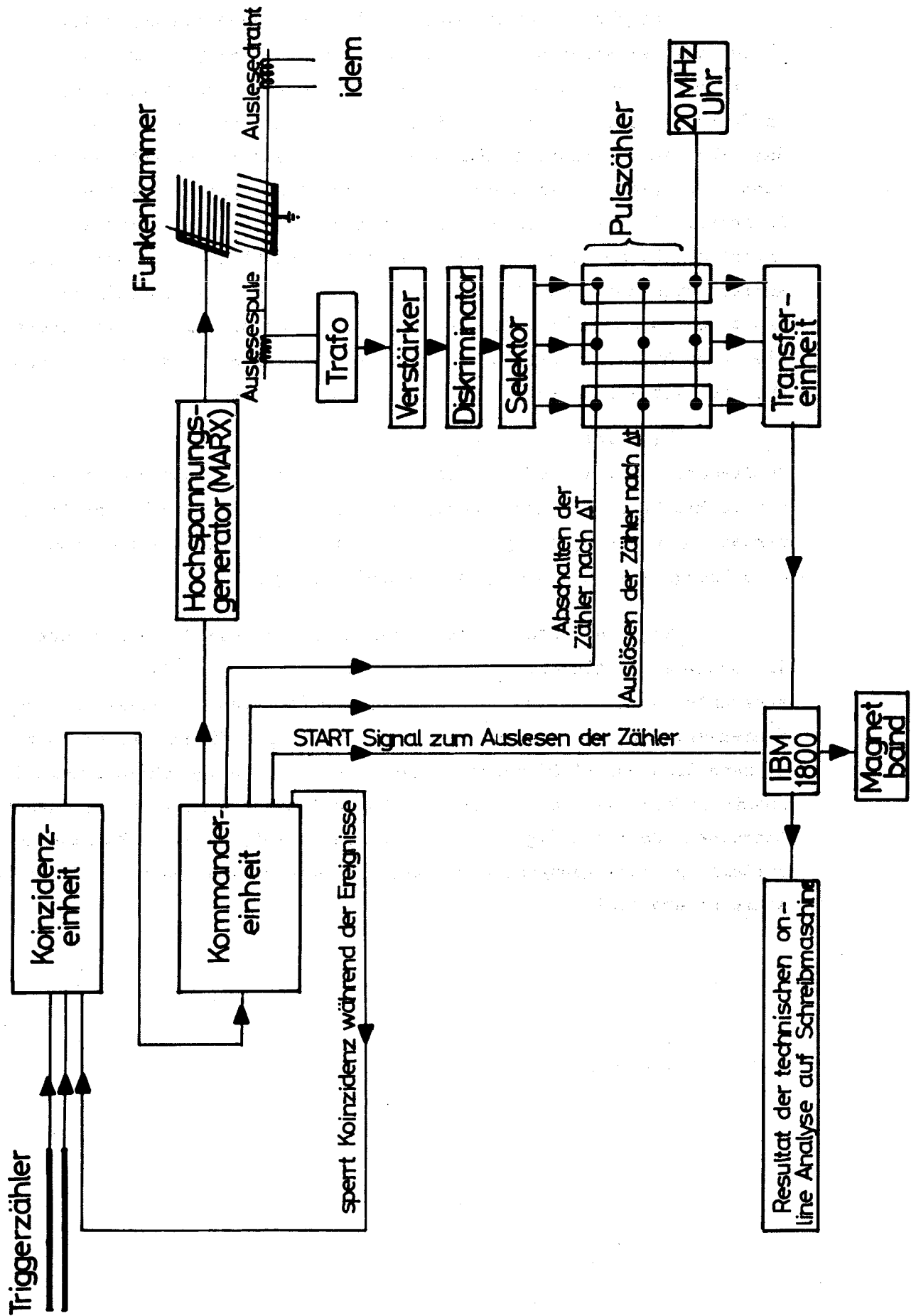
verstrichene Zeit in Sekunden an. Abgeschaltet wird ein Pulszähler durch ein Signal, das die Ankunft einer akustischen Welle in einer Auslesespule meldet. Dadurch ist die Laufzeit der akustischen Welle gemessen. Dieses Signal wird folgendermassen erhalten (die hier beschriebene Elektronik befindet sich hinter jeder Spule):

Das mit der akustischen Welle im Auslesedraht gekoppelte Magnetfeld erzeugt in der Auslesespule, wie im Abschnitt beschrieben, ein Signal von etwa $10 \mu\text{V}$. Durch einen Wirbelstromtransformator wird das Signal um einen Faktor 9 verstärkt. Eine hochspannungsfeste Isolierung zwischen Primär- und Sekundärwicklung erlaubt es, die Nulllinie der Signale auf Erdpotential zu bringen. Durch einen Vorverstärker und einen dreistufigen Endverstärker wird der positive Teil des Signales auf etwa 5 - 10 V verstärkt. Signale, deren Amplitude einen gewissen Schwellwert überschreitet, werden dann im Diskriminator differenziert. Der Nulldurchgang des differenzierten Signales löst jenen Puls aus, der über eine Selektoreinheit jeweils einen Pulszähler abschaltet. Die Aufgabe des Selektors ist es, bei mehreren magnetostriktiven Signalen für jedes Diskriminatorsignal einen anderen Pulszähler in zeitliche Reihenfolge auszuschalten. Dadurch können bei mehr als einem Funken in der Kammer die Laufzeiten so vieler akustischer Wellen (im magnetostriktiven Draht) gemessen werden, wie Pulszähler pro Auslesespule vorhanden sind. (Im Experiment wurden bis zu 8 Pulszähler pro Auslesespule verwendet).

Nach einem Zeitintervall ΔT , das länger als die Laufzeit der akustischen Welle von einem Ende des Auslesedrahtes bis zum anderen gewählt ist, gibt der Kommander ein Signal an die Pulszähler und schaltet jene ab, die noch nicht durch die magnetostriktiven Signale abgeschaltet wurden.

Anschliessend gibt der Kommander der Rechenmaschine (IBM 1800) den Befehl, über die Transfereinheit den Inhalt der Pulszähler auszulesen. Zwischen Eintreffen des Startsignales von der Koinzidenzeinheit und dem Zeitpunkt, zu dem die Rechenmaschine den letzten Pulszähler ausgelesen hat, und der Marxgenerator wieder aufgeladen ist, muss der Kommander ausserdem die Koinzidenzeinheit sperren, damit sich zwei Ereignisse nicht stören können.

Abb. 15: Blockschaltbild der Steuerung der Funkenkammern und der Ausleseelektronik



Aufgabe der Rechenmaschine ist es, neben dem Lesen der Zähler die gelesenen Daten auf Magnetband zu schreiben und das Funktionieren der Kammern und deren Auslesesystem zu überprüfen. Im Folgenden sollen kurz zwei Beispiele beschrieben werden, wie das Arbeiten der Kammern während des Experimentes überprüft werden kann. Da sich an jedem der beiden Enden des magnetostriktiven Auslesedrahtes eine Auslesespule befindet, erhält man, solange die Anzahl der Funken kleiner oder gleich der Anzahl der Pulszähler pro Auslesespule ist, doppelte Information über die Funkenkoordinate (Abstand des stromführenden Drahtes zur linken und zur rechten Spule). Dies wird bei der "on-line" Analyse der Funkenkammerdaten ausgenutzt. Abb. 16 zeigt das Ergebnis einer solchen Analyse.

Die Matrix im oberen Teil der Abb. 16 gibt für einen Auslesedraht an, wie oft die linke Spule i Funken und gleichzeitig die rechte Spule k Funken registriert hat. Damit kann überprüft werden, ob beide Spulen gleich gut arbeiten. Im Idealfall müssten alle Ereignisse in den Diagonalelementen liegen.

Wegen des konstanten Spulenabstandes muss für jeden Funken die Summe aus Abstände zur linken und zur rechten Spule einen konstanten Wert ergeben. Die Tabelle in Abbildung 16 zeigt eine solche Summenverteilung. Verschieben sich die Maxima dieser Verteilungen, so kann dies darauf hinweisen, dass sich die Schallgeschwindigkeit im Drahte ändert; verbreitern sich die Verteilungen, dann kann das darauf hinweisen, dass die Signalform so schlecht ist, dass das Maximum der Verteilung nicht genügend genau gefunden werden kann oder die Elektronik schlecht arbeitet.

Abb.16 On-line Überprüfung der Funkenkammern

COIL EFFICIENCY		M1 VERSUS		M2	→ k					
0	0	0	0	0	0	0	0	0	0	0
1	0	9	1	0	0	0	0	0	0	0
2	0	0	31	1	1	0	0	0	0	2
3	0	0	3	137	13	2	1	0	0	9
4	0	0	0	15	247	25	2	3	0	16
5	0	0	0	0	33	300	26	2	3	19
6	0	0	0	0	0	31	258	19	5	17
7	0	0	0	0	0	4	33	131	28	8
8	0	0	0	0	0	0	3	15	131	8

ASS-0.009 COW 10.5 82.1

Matrix, die für einen Auslesedraht angibt, wie oft für eine gewisse Anzahl von Ereignissen, die linke Spule i und gleichzeitig die rechte Spule k Funken registriert hat.

BEG 7700
BIN 5

98

1	1
2	2
3	2
4	0
5	0
6	1
7	1
8	6
9	1
10	3
11	8
12	10
13	2474
14	3255
15	7
16	0
17	0
18	0
19	0
20	0
21	0
22	0
23	0
24	0
25	0
26	0
27	0
28	0
29	0
30	0

Verteilung der Laufzeitsummen

Summe aus Laufzeit zur linken und laufzeit zur rechten
Spule in willkürlichen Einheiten 0.25µsec

IV. MESSUNGEN UND TESTRESULTATE

Um die Brauchbarkeit der Funkenkammern für ein Experiment zu prüfen, wurde Messgenauigkeit, Mehrteilchen Nachweiswahrscheinlichkeit und Gedächtniszeit der gebauten Kammern gemessen. Dazu wurden zwei Kammern mit jeweils 4 Elektroden hintereinander gestellt und durch Szintillationszähler verlangt, dass mindestens ein Teilchen die Kammer passiert.

1. Rekonstruktion der Spuren

Sei L der Abstand zwischen den Auslesepulen, v die akustische Geschwindigkeit im magnetostriktiven Draht, t_ℓ und t_r die im Pulszähler der linken bzw. rechten Spule gemessene Laufzeit der akustischen Welle und Δt das Zeitintervall zwischen Entladen der Kammer und dem Auslösen der Pulszähler (siehe Seite 29), dann gilt, wenn x der Abstand zwischen stromführendem Draht und linker Spule ist:

$$(t_\ell + \Delta t) = x/v$$

$$(t_r + \Delta t) = (L - x)/v$$

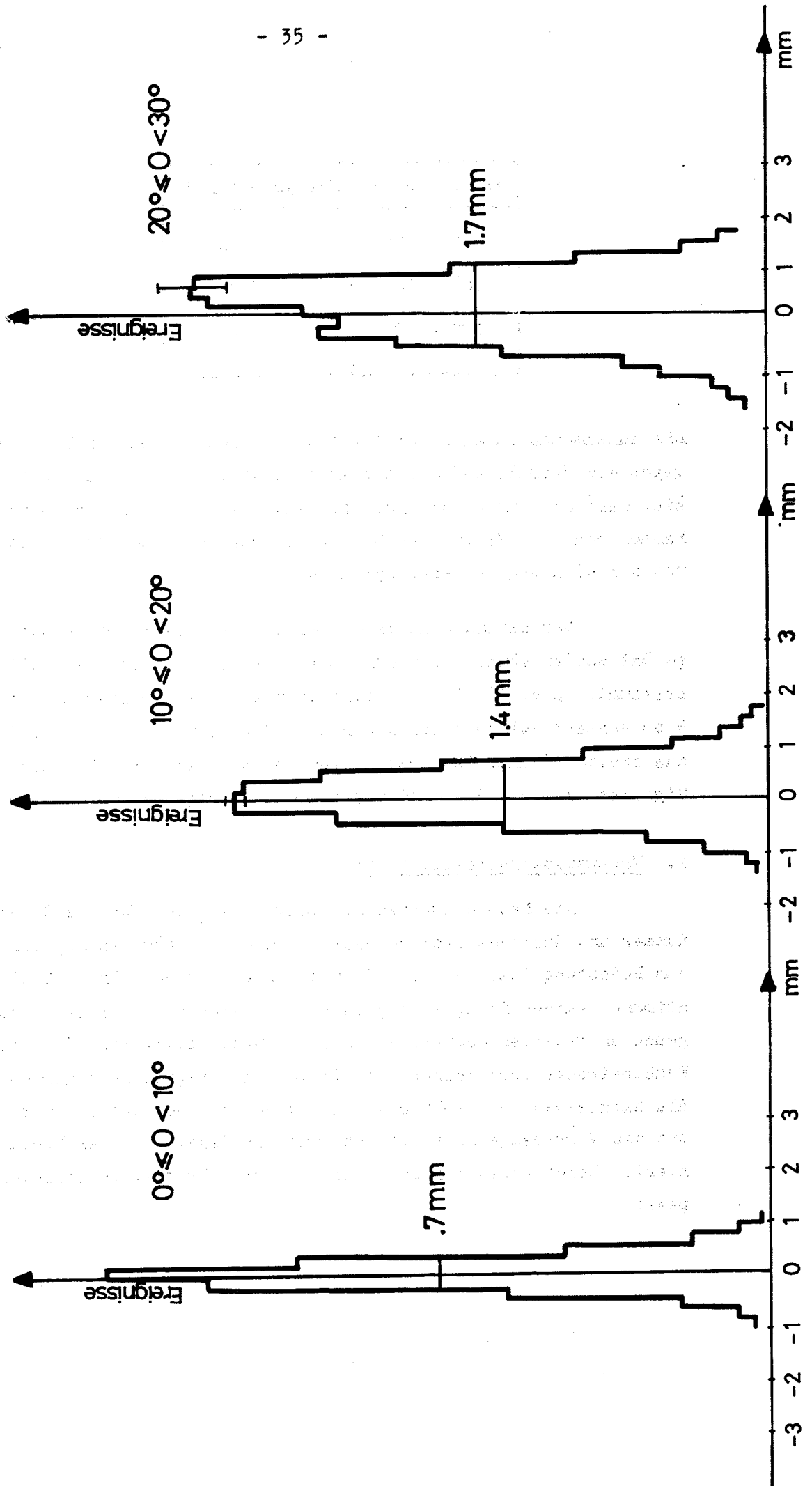
und damit unabhängig von Δt

$$x = L - \frac{(t_\ell - t_r)v}{2}$$

Somit sind die Koordinaten des Funkens im Bezug auf die Auslesepulen bekannt, nicht jedoch ihre absolute Lage im Raum. Um diese zu bestimmen wird senkrecht durch die Kammern ein gut kollimierter Teilchenstrahl mit bekannter Lage im Raum geschickt und der Koordinatenursprung in jeder Kammer als Schwerpunkt der gemessenen x -Werte definiert.

Eine Angabe über die Genauigkeit der Rekonstruktion des Funkenortes erhält man, indem man durch die in den beiden Kammern (d.h. 8 Drahtebenen) gemessenen 4 Koordinatenpaare Geraden legt und damit die wahrscheinlichste Spur definiert. Die Abweichungen der einzelnen Punkte von dieser Geraden sind dann ein Mass für die erreichte Genauigkeit der Rekonstruktion. Abb. 17 (Seite 35) zeigt solche Verteilungen für verschiedene Einfallswinkel der Teilchen. Aus Messungen an mehreren Kammern ergeben sich folgende Werte für den mittleren Fehler der Rekonstruktion der Spuren:

Abb. 17: Messgenauigkeit der Funkenkammern als Funktion der Neigung der Teichenbahn gegen die Feldrichtung in der Kammer (Differenz zwischen angepasster Gerader und gemessenem Punkt).



Winkelbereich	Messgenauigkeit
0° - 10°	± 0.4 mm
10° - 20°	± 0.7 mm
20° - 30°	± 1.0 mm

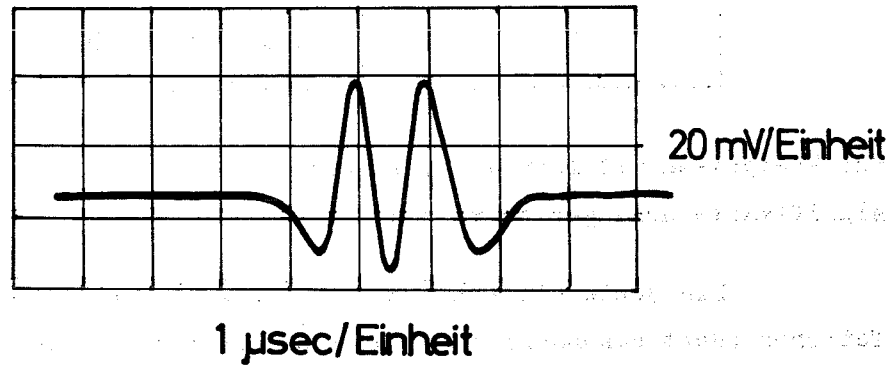
Die zunehmende Ungenauigkeit mit wachsendem Einfallswinkel der Spur gegen die Normale auf die Kammerebene stimmt mit der Beobachtung überein, dass die Funken mit grossem Winkel gegen die Feldrichtung in der Kammer nahe ($\sim 1/2$ cm) den Elektroden von der Teilchenspur abweichen, und der Richtung des angelegten Feldes folgen.

Der minimale Abstand, bei dem zwei Teilchen gerade noch aufgelöst werden können, wird durch die Breite des magnetostriktiven Signales bestimmt. Abb. 18 (S. 37) zeigt magnetostriktive Signale bei 2.5 und 5 mm Abstand zwischen den stromführenden Drähten. Bei 5 mm fällt gerade das zweite Minimum des ersten Signales auf das erste Minimum des zweiten Signales, so dass die beiden Signale gut aufgelöst werden können.

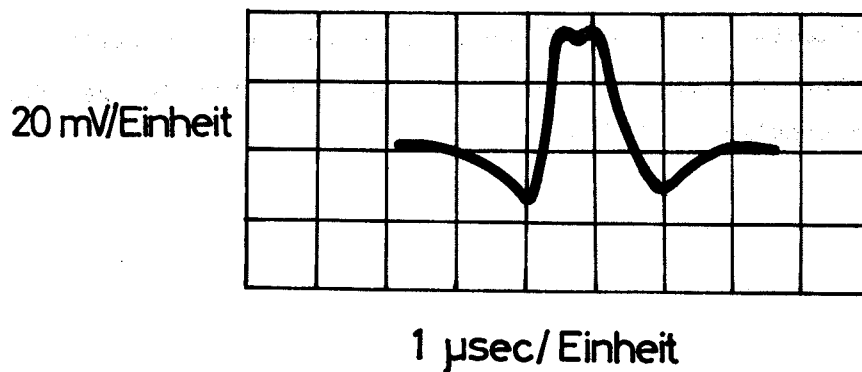
2. Nachweiswahrscheinlichkeit

Die Nachweiswahrscheinlichkeit ist nur für das Gesamtsystem Kammer und Auslesesystem gemessen worden, die für das Experiment allein von Bedeutung ist. Um sie für eine Funkenstrecke für n Teilchen zu bestimmen, wurden Ereignisse gesucht, in denen drei von den 4 Funkenstrecken genau n Teilchen registriert haben. Damit lässt sich für die vierte Funkenstrecke vorhersagen, wo die Teilchen registriert werden sollten. Als Nachweiswahrscheinlichkeit ist dann der Quotient aus der Zahl, wie oft die Vorhersage gestimmt hat, und der Gesamtzahl der Ereignisse definiert. Dabei ergaben sich folgende Werte für zwei verschiedene Ebenenpaare:

**Abb. 18: Auflösungsvermögen der magnetostriktiven
Auslesemethode (Signalabstand 5 mm)**



Signalabstand 2.5mm



Anzahl der Teilchen n	Nachweiswahrscheinlichkeit	
	<u>GAP 1</u>	<u>GAP 2</u>
1	98.2 ± 4%	97.6 ± 4%
2	97.2 ± 2%	95.0 ± 2%
3	93.4 ± 3.6%	92.4 ± 3.6%
4	93.0 ± 10%	88.3 ± 10%

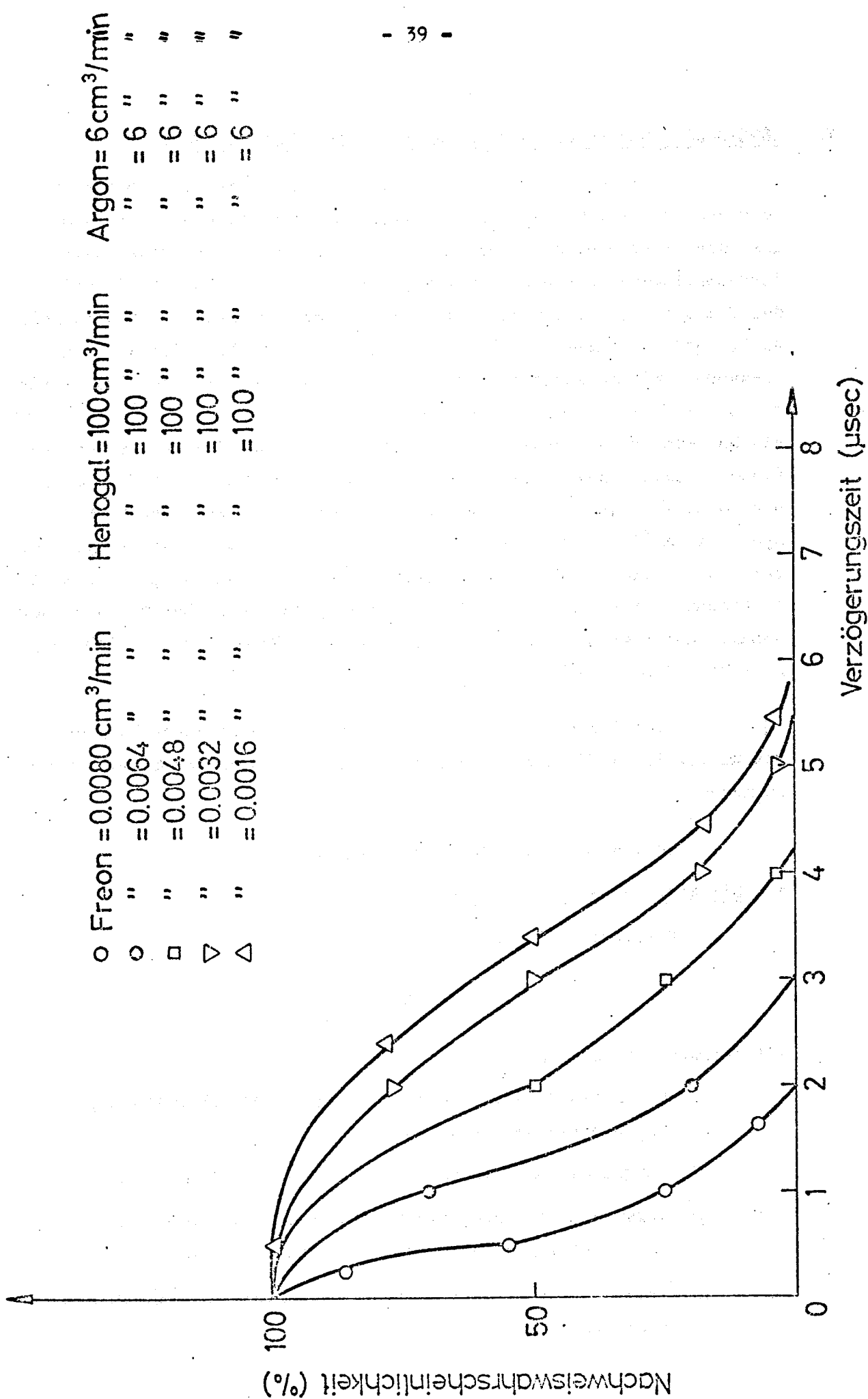
Für Ereignisse mit mehr als 4 Spuren war die Statistik zu klein, um signifikante Aussagen zu machen.

Die schlechtere Nachweiswahrscheinlichkeit für Ereignisse mit vielen Teilchen rührt besonders von den gegen das Feld in der Kammer geneigten Spuren her. Bei ihnen ist die Funkenlänge und damit der Widerstand gross. Dadurch wird der Funkenstrom so klein, dass das von ihm produzierte magnetostriktive Signal zu klein ist, um noch registriert zu werden.

3. Gedächtniszeit

Um die Gedächtniszeit bei verschiedenen Füllgasen der Kammern zu messen, wurde das Anlegen der Hochspannung (Verzögerungszeit t) verzögert und die Nachweiswahrscheinlichkeit der Kammer gemessen. Abb. 19 (Seite 39) zeigt Messungen für ein Gemisch von 90% Xenon und 10% Argon mit verschiedenen Freonbeimengungen. Während des Experiments liefen die Kammern mit einer Konzentration von 2.5×10^{-5} Freon, so dass der Wirkungsgrad bis zu einer Verzögerungszeit von 1.5 μsec praktisch 100% war.

Abb.19: Gedächtniszeit als Funktion des Gasgemisches



V. ANWENDUNG DER FUNKENKAMMERN IM CERN-BOSONSPEKTROMETER

Im Folgenden soll das Experiment: CERN-Bosonspektrometer beschrieben werden, bei dem die Drahtfunkenkammern eingesetzt wurden. Ziel des Experimentes war es, das Massenspektrum der nichtseltsamen, einfachgeladenen Bosonen in der Reaktion $\pi p \rightarrow pX$ (X ist hierbei das Teilchensystem, das neben dem Proton erzeugt wird) im Massenbereich um 1.3 GeV zu studieren. In diesem Massenbereich befindet sich das A_2 -Meson. Wie in Kapitel V, Seite 51 genauer beschrieben wird, wurde es zuerst¹⁹⁾ in Blasenkammerexperimenten in der Reaktion $\pi^+ p \rightarrow p(\pi\pi\pi)^+$ als Maximum in der Verteilung der effektiven Masse des 3π -Systems gefunden. Seine Masse wurde zu $M = 1305$ MeV bestimmt, seine Halbwertsbreite zu $\Gamma = 90$ MeV. In einem Zählerexperiment (CERN Missing Mass Spektrometer)²⁰⁾ wurden mit viel Statistik und gutem Auflösungsvermögen bei 6 GeV/c und 7 GeV/c Eingangsimpuls des π -Mesons an der Stelle des A_2 -Mesons zwei schmale Verteilungen von ungefähr gleichem Wirkungsquerschnitt und gleicher Breite gefunden ($M_1 = 1.275$ GeV, $M_2 = 1.318$ GeV, $\Gamma_1 = \Gamma_2 = 0.021$ GeV).

Ziel dieses Experimentes ist es, die Massenverteilung des A_2 -Mesons bei 2.6 GeV/c Eingangsenergie, 400 MeV/c über dem Schwellenimpuls zu messen.

Dazu wurde folgende Methode angewandt:

1. Die Methode der fehlenden Masse²¹⁾

Betrachten wir die Reaktion

$$(1) + (2) \rightarrow (3) + (4) \quad (a)$$

mit folgender Bedeutung der Symbole

- 1) bekanntes Strahlteilchen mit gemessenem Laborimpuls p_1 und Laborgesamtenergie E_1 .
- 2) Targetteilchen, im Laborsystem ruhend, mit Masse m_2 .
- 3) Irgendein Teilchen, dessen Impuls p_3 und Winkel Θ_3 im Laborsystem gemessen und dessen Masse bestimmt wird.
- 4) Irgendein System von Teilchen mit unbekannter effektiver Masse M_4 .

Die Methode der fehlenden Masse besteht darin, Impuls und Streuwinkel des Teilchens (3) zu messen und daraus, unter Verwendung des Energie- und Impulssatzes, die effektive Masse des Systemes (4) zu bestimmen. Als "fehlende Masse" des Teilchens (3) (MM_3) wird die neben dem Teilchen noch vorhandene Masse, also die effektive Masse des Systems (4) bezeichnet.

$$MM_3 = M_4 = (E_4^2 - p_4^2)^{1/2} = [(E - E_3)^2 - (\vec{p}_1 - \vec{p}_3)^2]^{1/2}$$

wobei E die Gesamtenergie im Laborsystem bezeichnet

$$E = m_2 + E_1$$

Damit erhält man:

$$MM_3^2 = M_4^2 = (E_1 + m_2 - E_3)^2 - p_1^2 - p_3^2 + 2p_1 p_3 \cos \Theta_3 \quad (b)$$

Resonanzzustände des Systems (4) zeigen sich in einem M_4 -Spektrum als Maxima über einem Untergrund aller möglichen Reaktionen, die durch (a) beschrieben werden können.

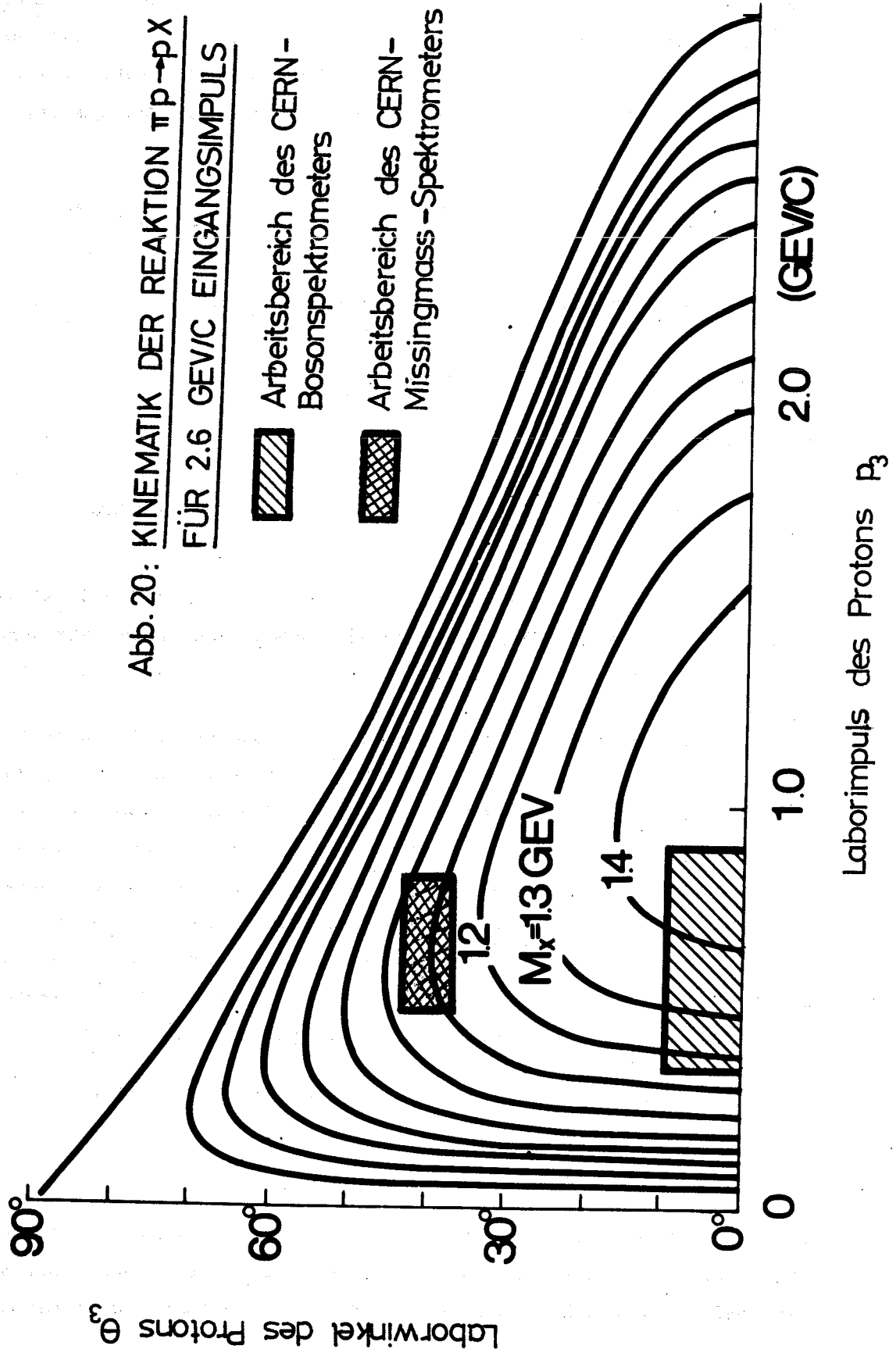
Diese Methode, nach Resonanzen zu suchen, hat den Vorteil, prinzipiell einfach und, da sie sehr viele Reaktionskanäle erfasst, sehr allgemein zu sein. Ihr Nachteil ist, zumindest wenn man sich auf die Messung von p_3 und Θ_3 beschränkt, dass sie Resonanzzustände nur über einem sehr grossen Untergrund zeigt und nur wenige Aussagen über deren Quantenzahlen zulässt. Speziell für die untersuchte Reaktion $\pi^- p \rightarrow p X^-$ lässt sich über die Quantenzahlen des Systems X folgendes aussagen: Baryonenzahl $B = 0$, Seltsamkeit $S = 0$, Isotopenspin $I = 1$ oder 2 , dritte Komponente des Isotopenspins $I_3 = -1$.

In Abbildung 20 (Seite 42) sind für die Reaktion



für $E_1 = 2.6$ GeV die nach (b) berechneten Linien konstanter fehlender Masse MM_3 als Funktion von p_3 und Θ_3 bezeichnet. Aus ihnen sieht man, dass bei gegebenem Eingangsimpuls dann die höchste Masse erreicht werden

Abb. 20: KINEMATIK DER REAKTION $\pi p \rightarrow pX$
FÜR 2.6 GEV/C EINGANGSIMPULS



kann, wenn das Teilchen (3) unter 0° gemessen wird. Im beschriebenen Experiment wurden in der Reaktion $\pi p \rightarrow pX$ die Protonen (3) im schraffierten Bereich von Abb. 20 gemessen.

Formel (b) enthält als Winkelabhängigkeit einen $\cos\Theta_3$ -Term, so dass bei 0° M_4 unabhängig von Θ_3 ist. Deshalb braucht zur Bestimmung von M_4 der Winkel Θ_3 nur ungenau bekannt zu sein, dagegen muss p_3 genau gemessen werden.

Ein weiterer Vorteil der Teilchenproduktion unter 0° ist die Tatsache, dass bei hohen Energien die meisten Resonanzen durch Austauschreaktionen erzeugt werden, bei denen die Wirkungsquerschnitte um 0° maximal sind. Andererseits gibt es auch Reaktionen, bei denen Teilchen ausgetauscht werden sollten, die noch nicht experimentell gefunden wurden und deren Existenz damit sehr fraglich ist. Bei diesen Reaktionen sind die Wirkungsquerschnitte bei 0° daher sehr klein. Spezielle Beispiele dafür sind der Austausch von Objekten mit Isospin $I = 2$, Seltsamkeit $S = 2$ und gewisse Kanäle von Vektormesonenaustausch²²).

2. Experimenteller Aufbau, Messmethode und Abschätzung der Messfehler

Das Experiment (1. Phase des CERN-Bosonspektrometers CBS) wurde 1967-1968 im Teilchenstrahl q_5 am CERN-Protonensynchrotron durchgeföhrt²³). Der q_5 Strahl ist ein niederenergetischer (zwischen 0.5 GeV/c und 3.0 GeV/c), unseparierter Teilchenstrahl, der von einem maschineninternen Berylliumtarget (20 mm lang und 1 mm² Querschnitt) unter 350 mrad Produktionswinkel erzeugt wird. Die Strahloptik blendet aus dem Produktionsspektrum des Be-targets einen Winkelbereich von 0.5×10^{-4} sterad und ein Impulsband von 3% aus. Der Zentralwert des Strahles wurde während eines Eichruns mit dem unten beschriebenen Spektrometer auf $\pm 0.3\%$ genau gemessen. Die Abweichung des Impulses der einzelnen Teilchen vom Zentralimpuls wurde mit Hilfe von drei Szintillationszählerhodoskopen und den Strahlmagneten ebenfalls auf $\pm 0.3\%$ genau ermittelt. Die Art der einfallenden Teilchen wurde durch zwei Schwellenwert-Cerenkovzähler identifiziert. Bei einem Impuls von 2.60 GeV/c bestand der Strahl je nach Vorzeichen der Strahlmagneten aus:

28.3% π^+ , 0.9% K^+ , 70.8% p und etwa 1% d

bzw. aus:

98.1% π^- , 1.2% K^- und 0.7% p^- .

Abbildung 21 (Seite 45) zeigt eine schematische Skizze des experimentellen Aufbaus: zwei Szintillationszählerrhodoskope H1 und H2 bestimmen auf ± 1 mrad genau die Richtung, mit welcher das einfallende Teilchen auf ein Target aus flüssigem Wasserstoff auftrifft. Strahlabwärts hinter dem Target befindet sich ein ein-armiges Magnetspektrometer. Dieses besteht aus einem Paar von Drahtfunkenkammern (SC1 und SC2) vor und einem weiteren Paar (SC3 und SC4) hinter einem grossen H-Magneten (1.5 m $\times$ 0.5 m Öffnung) und einem grossen Flugzeitszintillationszähler (R), der die empfindliche Fläche der Funkenkammer SC4 überspannt. Die Funkenkammern SC3, SC4 und der Zähler R sind gemeinsam auf einem Drehtisch montiert, so dass sie verschiedene Ablenkwinkelbereiche überspannen können. Mit Hilfe der Funkenkammerpaare SC1 - SC2 und SC3-SC4 wird die Ablenkung der Rückstossprotonen der Reaktion $\pi p \rightarrow pX$ durch das Magnetfeld des H-Magneten gemessen und daraus deren Impuls bestimmt.

Zwischen dem Szintillationszähler T2 unmittelbar vor dem Wasserstofftarget und dem grossen Szintillationszähler R wird die Flugzeit (T.O.F = Time-of-flight) der Rückstossprotonen gemessen und daraus ihre Geschwindigkeit β bestimmt. Die beiden Kammern SC1-SC2 dienen ausserdem dazu, die Richtungen der geladenen Zerfallsprodukte von X zu bestimmen und damit den Vertex der Reaktion zu finden. Die Zahl der geladenen Zerfallsprodukte von X wird durch die Kammern SC1, SC2 und die 4 Szintillationszähler (D1, D2, D3, D4, in der Zeichnung (D) sind nur zwei eingezeichnet) bestimmt, die zusammen praktisch den gesamten Raumwinkel um das Target überdecken.

Ein Ereignis wird logisch definiert durch:

- die beiden Szintillationszähler T1 und T2, die das einfallende Teilchen feststellen.
- Den Szintillationszähler R, der in einem bestimmten Zeitintervall, nachdem von T2 das einfallende Teilchen registriert wurde, ein geladenes Teilchen verlangt. Das Zeitintervall ist so gewählt, dass langsame Protonen zwischen 300 und 800 MeV/c akzeptiert werden, nicht jedoch π -Mesonen, deren Geschwindigkeit wegen ihrer kleinen Ruhemasse nahe der Lichtgeschwindigkeit liegt.

- Einen der beiden Szintillationszähler V1 oder V2. Sie befinden sich hinter der Funkenkammer SC2. Ihre Fläche ist, bis auf ein Loch in deren Mitte (damit Rückstossprotonen ohne ein Signal zu geben in den Magneten gelangen können), gleich der Fläche der Funkenkammer SC2. Sie verlangen mindestens ein geladenes Zerfallsprodukt des Systems X.
- Den Zähler B, der nicht ansprechen darf und damit feststellt, ob das einfallende Teilchen aus dem Strahl gestreut wurde.

Die Bedingung für ein Ereignis ist also: $T1 \cdot T2 \cdot \bar{B} \cdot V1/V2 \cdot R$.

Der Impuls des Protons wird aus dessen Ablenkung im Magnetfeld gemessen. Die Genauigkeit der Impulsmessung wird durch die Vielfachstreuung in den Drähten und Fenstern der Drahfunkenkammern begrenzt. Das durchquerte Material entspricht etwa 0.16 Strahlungslängen, was für Protonen von 500 MeV/c eine Winkelungenauigkeit $\Delta\delta$ von ± 6.8 mrad bringt. Ihre Ablenkung δ im Magnetfeld beträgt bei dem meist eingestellten Feld von 4.15 kGauss etwa 700 mrad, so dass sich eine Impulsauflösung der magnetischen Analyse $\frac{\Delta\delta}{\delta} = \frac{\Delta p_{\text{magnet}}}{p_{\text{magnet}}}$ von 0.97% ergibt. Zwischen den Zählern T2 und R (ihr Abstand ist etwa 6 m) wird die Flugzeit der Teilchen auf ± 0.3 nsec mit einem relativen Fehler von 0.7% bestimmt. Daraus lässt sich die Geschwindigkeit β der Teilchen berechnen. Impuls und Geschwindigkeit erlauben es, die Masse der Teilchen zu berechnen und die Protonen zu identifizieren. Abb. 22 (Seite 47) zeigt eine solche Massenverteilung. Position und Breite dieser Verteilung erlauben es, die Genauigkeit und Stabilität des Systems zu überwachen. Ist das Teilchen als Proton identifiziert, lässt sich aus der Geschwindigkeit der Impuls berechnen. Damit ergibt sich für die Impulsauflösung aus der Flugzeitmessung:

$$\frac{\Delta p(\text{TOF})}{p(\text{TOF})} = 0.9\% \text{ für } p = 500 \text{ MeV/c.}$$

Kombiniert man die beiden mit ihren Fehlern gewichteten Messwerte für den Protonimpuls, so lässt sich für den Impuls des Protons beim Verlassen des Wasserstofftargets eine Messgenauigkeit von $\frac{\Delta p}{p} = 0.66\%$ erreichen. Um den Impuls p_z des Protons am Vertex zu erhalten, muss noch der Energieverlust der Protonen im Wasserstofftarget berücksichtigt werden. Der Vertex lässt sich als Schnittpunkt der einfallenden Teilchenspur mit den Richtungen der Zerfallsprodukte - gemessen in den ersten beiden Kammern - auf ± 1.5 cm in Strahlrichtung bestimmen. Das ergibt einen

•• MASS OF PROTON ••

RUN #120304

MASS SPECTRUM
OF RECOIL PARTICLES
IN T-P. (2.65 GeV/c)

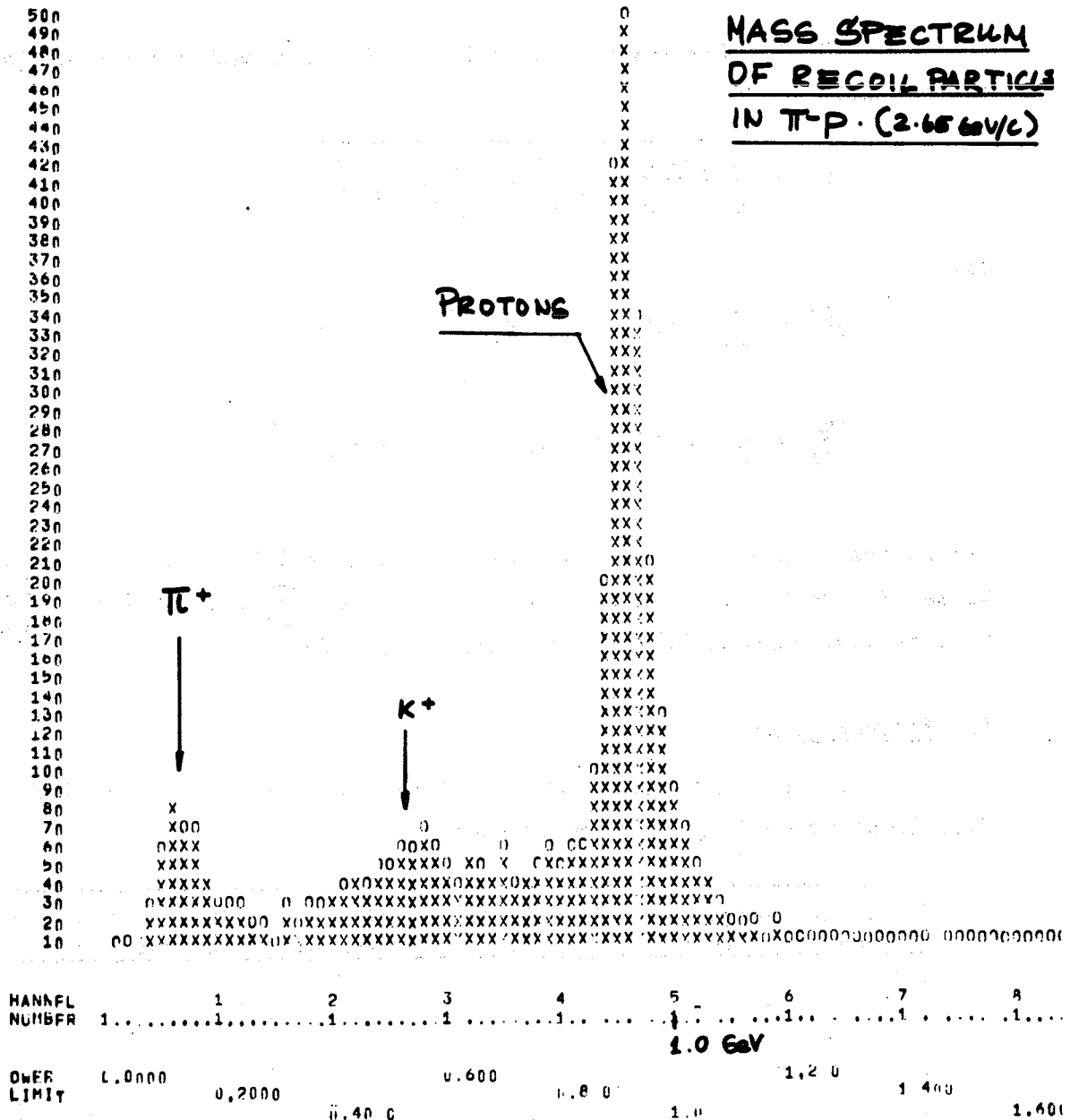


Abb. 22: Massenverteilung der vom Magnetspektrometer (CBS) akzeptierten Teilchen. Für das Mesonenspektrum werden nur die Rückstossprotonen ($840 < \text{Masse} < 1040 \text{ MeV}$) genommen. Typisches Beispiel eines Produktionsruns.

Beitrag zur Impulsungenauigkeit von

$$\frac{\Delta p(\text{Vertex})}{p} = 0.62\%,$$

was zusammen mit obigen 0.66% eine Impulsunschärfe am Wechselwirkungs-
ort von 0.9% ergibt.

Für die Massenauflösung bei $\Theta_3 = 0^\circ$ gilt nach der Formel (b)
im Abschnitt V

$$\Delta M_4 = \sqrt{\left(\frac{\partial M_4}{\partial p_1} \Delta p_1\right)^2 + \left(\frac{\partial M_4}{\partial p_3} \Delta p_3\right)^2}$$

mit

$$\frac{\partial M_4}{\partial p_1} = \frac{1}{M_4} \left\{ (E_1 + m_2 - E_3) \beta_1 - p_1 + p_3 \right\}$$

$$\frac{\partial M_4}{\partial p_3} = \frac{1}{M_4} \left\{ (E_3 - E_1 - m_2) \beta_3 - p_3 + p_1 \right\}$$

Daraus ergibt sich für $p_3 = 500 \text{ MeV/c}$ - d.h. bei $E_1 = 2.6 \text{ GeV}$ bei einer
Masse M_X von 1.3 GeV - mit obigen Werten für $\frac{\Delta p_3}{p_3} = \pm 0.9\%$ und $\frac{\Delta p_1}{p_1} =$
 $\pm 0.3\%$ eine Massenauflösung von $\Delta M_4 = \pm 5.4 \text{ MeV}$.

3. Messresultate

In der erste Phase (November 1967 - Januar 1968), lief das
Experiment unter folgenden Bedingungen:

Eingangsenergie E_1 des π^- in GeV	Magnetfeld in KG	Winkel des Drehtisches gegen die Strahlrichtung
2.60	3.0	27°
2.60	2.0	18°
2.55	2.0	18°
2.65	2.0	18°

Eingangsenergie, Magnetfeld und Winkel des Drehtisches gegen die Strahlrichtung wurden geändert, um systematische Fehler der Apparatur aufzufinden, die Strukturen im Massenspektrum vortäuschen könnten. So bewirkt eine Änderung des Eingangsimpulses von 2.55 nach 2.65 GeV/c bei festgehaltenem Rückstossimpuls des Protons im Laborsystem eine Verschiebung von M_X um 30 MeV. Strukturen, die dadurch entstünden, dass die Apparatur gewisse Impulse bevorzugte, würden deshalb bei den beiden verschiedenen Eingangsenergien im Massenspektrum um 30 MeV verschoben erscheinen und leicht erkannt werden.

Abb. 23 zeigt das in dieser ersten Phase gemessene Massenspektrum. Für das Histogramm wurden nur Ereignisse ausgewählt, bei denen der Winkel des Rückstossprotons im Schwerpunktsystem zwischen 176 und 180 Grad liegt, um für das System X zwischen 1200 MeV und 1400 MeV eine konstante geometrische Akzeptanz zu erhalten. Ausserdem wurde verlangt, dass neben dem Rückstossproton mindestens drei weitere geladene Teilchen in der Reaktion erzeugt wurden. Da das A_2 mit etwa 85% Wahrscheinlichkeit über $\rho \pi$ in 3π -Mesonen zerfällt, war es damit möglich, die statistische Signifikanz der Struktur zu erhöhen, ohne die Struktur selbst zu verändern.

Um bei höherem Magnetfeld arbeiten zu können, die geometrische Akzeptanz des Systems zu erhöhen und mögliche Fehler der Apparatur zu finden, wurde die Geometrie des Spektrometers für die zweite Phase des Experimentes (Februar bis Mai 1968) geändert. In ihr lief das Experiment unter folgenden Bedingungen:

Eingangsimpuls	Magnetfeld	Winkel Drehtisch-Strahl
π^+ 2.65 GeV/c	5.5 KG	30°
π^- 2.65 GeV/c	4.15 KG	24°

Abb. 23b), c) zeigen die gemessenen Massenspektren wiederum für drei oder mehr geladene Zerfallsprodukte und für Winkel des Rückstossprotons im Schwerpunktsystem von mehr als 176°. In b) und c) ist die Masse M_X allein aus der Flugzeitmessung des Rückstossprotons berechnet. Abb. 24a zeigt die gesamten zwischen November 1967 und Mai 1968 gemessenen Daten. Darunter (Abb. 24b) sind als Vergleich die vom "CERN-Missing-Mass-Spektrometer" 1965-1967 bei Eingangsenergien von 6.0 und 7.0 GeV gemessenen Daten²⁰⁾ gezeichnet.

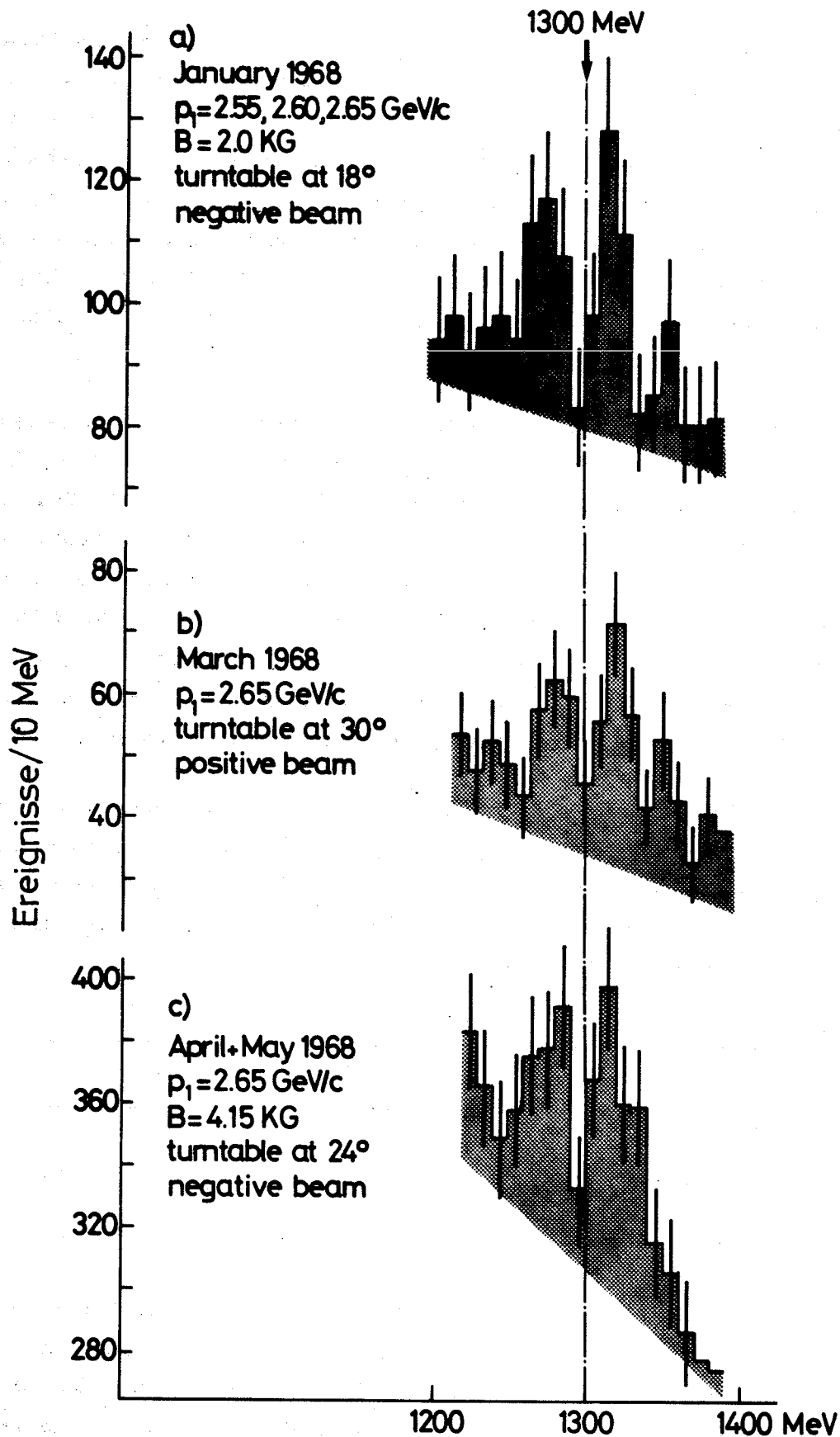


Abb.23: MASSENSPEKTREN DER REAKTION $\pi^- p \rightarrow p X^-$
ERHALTEN MIT DEM CERN BOSONSPEKTROMETER
(1968) UNTER VERSCHIEDENEN EXPERIMENTELLEN
BEDINGUNGEN.

Sie wurden aus der Analyse des Rückstossprotons in der Reaktion $\pi^- p \rightarrow p X^-$ im Jakobischen Winkel - dem Maximalwinkel im Laborsystem bei konstanter fehlender Masse (kariierter Bereich in Abb. 20, Seite 42) erhalten. Beide Spektren zeigen, abgesehen vom Untergrund, innerhalb der statistischen Fehler die gleiche Struktur, nämlich zwei symmetrische Verteilungen mit den Maxima bei:

$$M_X = M_1 = 1278 \pm 5 \text{ MeV}$$

$$M_X = M_2 = 1318 \pm 5 \text{ MeV}$$

und Halbwertsbreiten von

$$\Gamma_1 = \Gamma_2 = 22 \pm 5 \text{ MeV}.$$

Abbildung 24c zeigt die Summe von 24a und 24b. Dieselbe Verteilung ist in Abb. 25 gezeichnet, jedoch mit einer Intervalleinteilung von 5 MeV.

4. Vergleich der Messresultate mit Ergebnissen anderer Experimente und mögliche Interpretationen der gemessenen Strukturen

Zum ersten Male wurde von G. Goldhaber et al.¹⁹⁾ eine breite Struktur (zwischen 1 und 1.4 GeV) in der effektiven Masse des $\rho\pi$ -Systems in der Reaktion $\pi^+ p \rightarrow p \rho^0 \pi^+$ bei 3.65 GeV Eingangsenergie gemessen. Kurze Zeit danach wurde diese breite Struktur in zwei getrennte Systeme aufgelöst und als zwei verschiedene Resonanzen A1(1080) und A2(M = $1305 \pm 5 \text{ MeV}$, $\Gamma = 90 \pm 10 \text{ MeV}$) interpretiert²⁴⁾.

a) Eine einzelne Resonanz

Der Resonanz A₂ wurden folgende Quantenzahlen zugeschrieben:

- Der starke Zerfall in 3π -Mesonen verlangt negative G-Parität.
- Die Beobachtung von geladenen Zuständen verlangt einen Isospin I gleich oder grösser als 1. Hätte das A₂-Meson den Isospin I = 2, dann würde sich, wegen Erhaltung des Isospins bei der Produktion $\sigma(\pi^- n \rightarrow A_2 p) : \sigma(\pi^- n \rightarrow A_2 n) = 4 : 1$ ergeben. Gefunden wurde jedoch ungefähr 0, sodass dem A₂ der Isospin 1 zugeordnet wurde.
- Da das A₂ nur über einem grossen Untergrund von $\rho\pi$ erzeugt wird, gelang

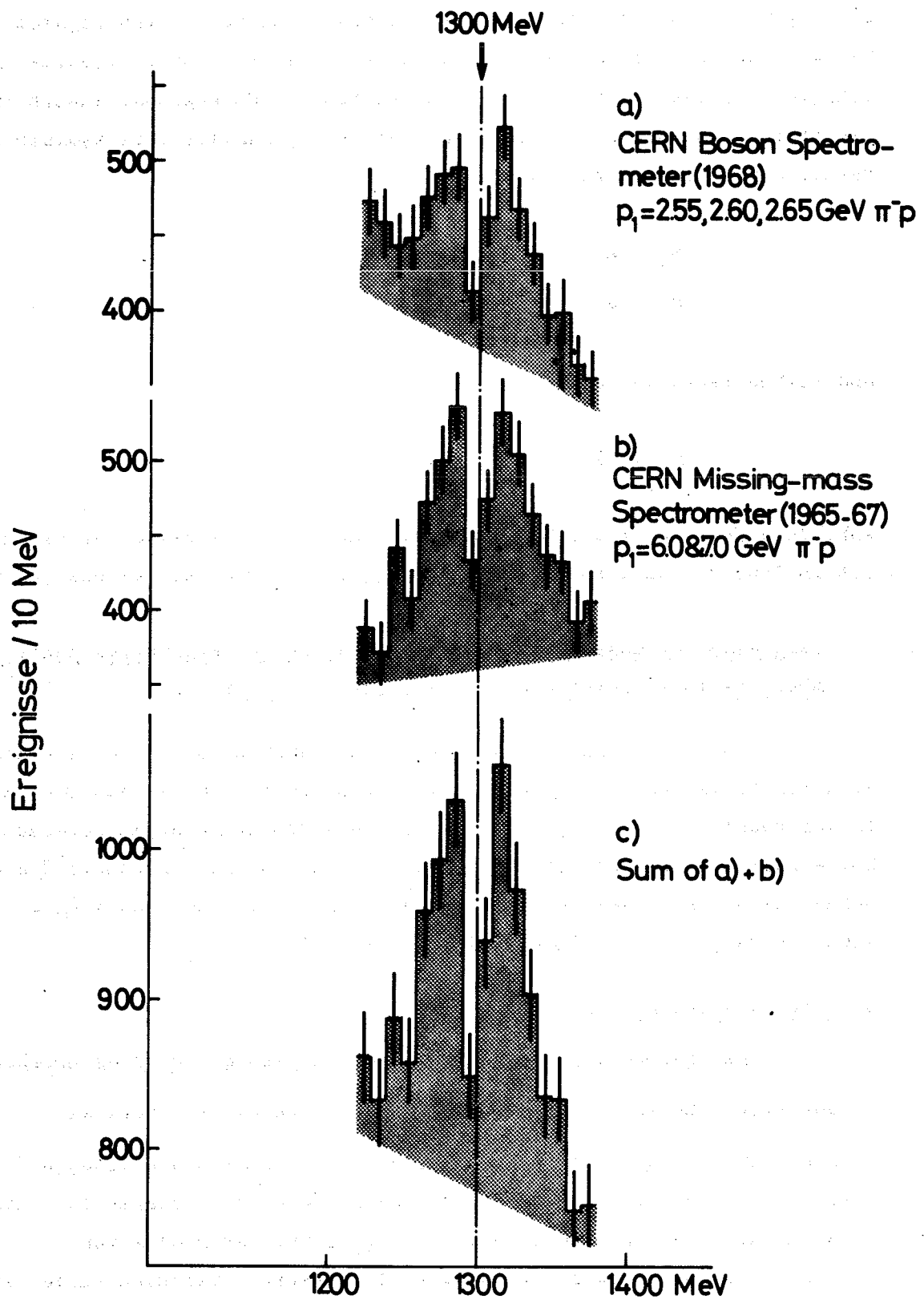


Abb. 24: SUMME DER GESAMTEN A_2 -DATEN VOM
CERN MISSING MASS EXPERIMENT 1965-67 +
CERN BOSONSPEKTROMETER 1968.

keine schlüssige Bestimmung von Drehimpuls und Parität J^P aus dem $\rho\pi$ Zerfall. 2^+ wird bei bestimmten Annahmen über den Untergrund bevorzugt, nur 0^+ ist wegen des Zerfalls in $\rho\pi$ ausgeschlossen.

Die Suche nach anderen Zerfallsarten hat ergeben²⁵⁾:

$A_2 \rightarrow \pi\rho$	$K\bar{K}$	$\pi\eta$	$\eta'\pi, (\pi\pi\pi) \neq \rho\pi$
$85.1 \pm 2.2\%$	$4.0 \pm 1.1\%$	$10.9 \pm 2.1\%$	?

Aus diesen Zerfallsarten kann J^P folgendermassen bestimmt werden:
Da sowohl η als auch π pseudoskalare Teilchen sind, folgt aus dem $\eta\pi$ -Zerfall, dass $J^P(A_2) = 0^+, 1^-, 2^+, 3^- \dots$ sein muss. Die G-Parität des $K\bar{K}$ -Systems berechnet sich aus: $G(K\bar{K}) = (-1)^{I+L}$ (L: relativer Bahndrehimpuls der beiden K's). Da $G(A_2) = -1$ und $I(A_2) = 1$ ist folgt mit G- und I-Erhaltung aus dem $K\bar{K}$ -Zerfall ($L = J$, da K spinlos): $J^P(A_2) = 0^+, 2^+, 4^+$. Wegen des $\rho\pi$ -Zerfalls ist 0^+ ausgeschlossen, und es folgt - unter der Annahme dass all diese Zerfälle einem einzigen Teilchen zuzuschreiben sind: $J^P(A_2) = 2^+(4^+)$, $I^G = 1^-$.

Wegen dieser Quantenzahlen wurde das A_2 zusammen mit den Teilchen $f^0(1260)$, $f^0(1514)$, und $K^*(1420)$ in das 2^+ -Nonet des SU_3 -Ordnungsschemas eingeordnet. Sowohl was Masse als auch Partialbreiten anbelangt, passt es gut in diese Ordnungsschema.

Die Hypothese einer einzigen Resonanz mit Breit-Wigner Form erscheint jedoch unwahrscheinlich, wenn man die Massenverteilungen von Abb.24 betrachtet. Eine Anpassung einer Breit-Wigner-Kurve an das gemessene Spektrum ergibt nur eine χ^2 -Wahrscheinlichkeit $P(\chi^2)$ von $< 10^{-4}$. Dies wurde auch von einem BNL-Blasenkammerexperiment (BNL = Brookhaven National Laboratory)²⁶⁾ bestätigt, wenn auch mit geringerer statischer Signifikanz.

b) Zwei Resonanzen mit verschiedenen Quantenzahlen

Eine grössere Wahrscheinlichkeit $P(\chi^2)$ - jeweils 20% für die getrennten Spektren Abb. 24 a, b (Seite 52) und 0.2% für deren Summe (Abb. 25, Seite 54) - ergibt die Hypothese zweier nichtinterferierender Resonanzen. Dabei ergeben sich aus einer Anpassung mit zwei inkohärenten Breit-Wigner-Kurven die Werte:

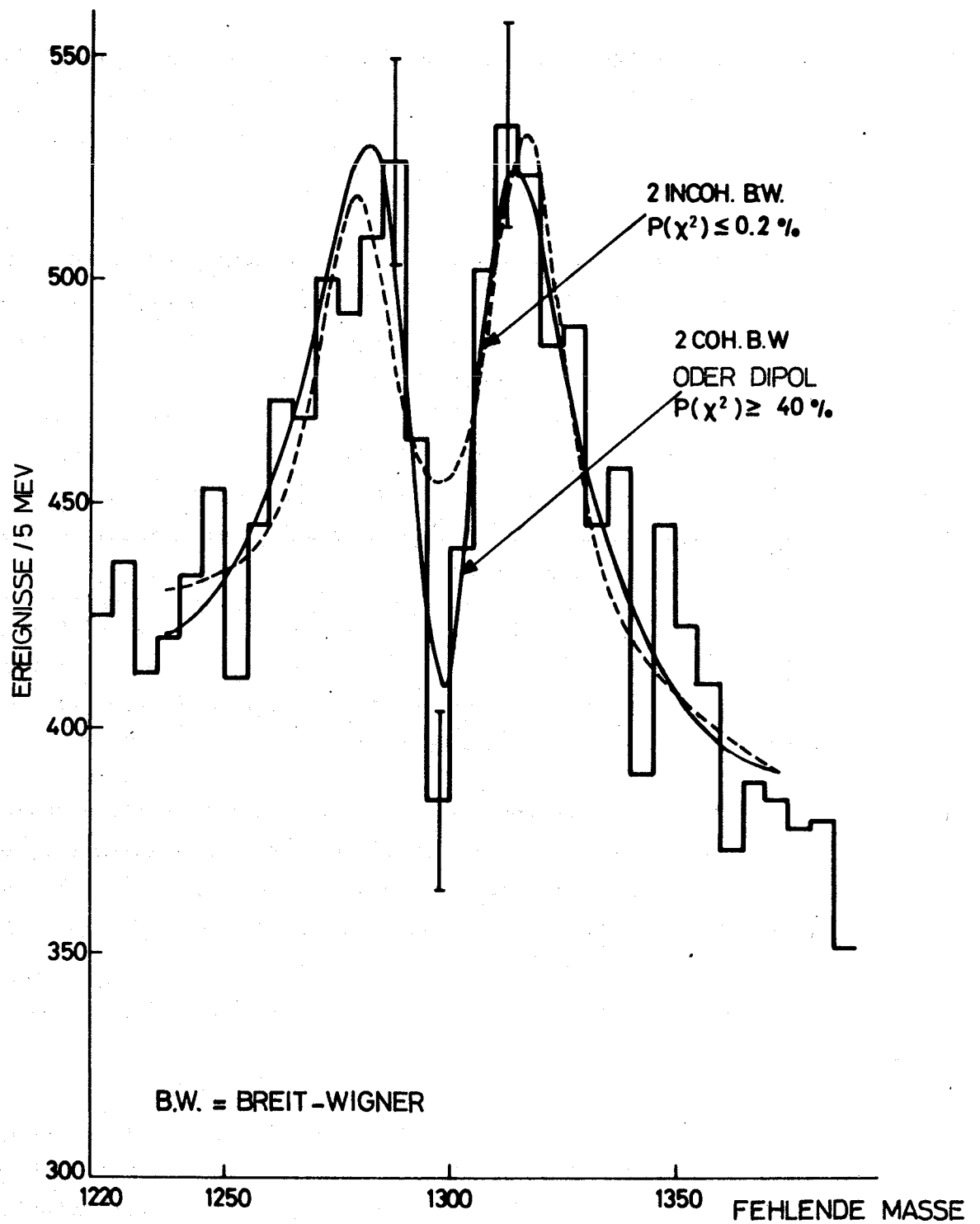


Abb.25: GESAMTES MASSENSPEKTRUM (CERN MISSING -
MASS EXPERIMENT + CERN BOSONSPEKTROMETER)
GEFITTET MIT VERSCHIEDENEN HYPOTHESEN.

$$A_2^{\text{low}} : M_1 = 1278 \pm 5 \text{ MeV} \quad \Gamma^{\text{low}} = 22 \pm 5 \text{ MeV}$$

$$A_2^{\text{high}} : M_2 = 1318 \pm 5 \text{ MeV} \quad \Gamma^{\text{high}} = 21 \pm 5 \text{ MeV}$$

Die Ergebnisse des BNL-Experimentes²⁶⁾, das in den Reaktionen $\pi^- p \rightarrow n K_1^0 K_1^0$ und $\pi^- p \rightarrow n K_1^0 K_1^0 + \text{neutrale Teilchen}$ nur ein Teilchen mit der Masse 1311 MeV und einer Breite $\Gamma = 21$ MeV misst, während es in $\pi^- p \rightarrow p(\rho\pi)^-$ zwei Maxima mit M_1 und M_2 sieht, und die Tatsache, dass die Summe aller Blasenkammerdaten der Reaktionen $\pi^- p \rightarrow n K_1^0 K_1^0$ und $\pi^\pm p \rightarrow p K_1^0 K^\pm$ ein Maximum bei etwa 1320 MeV ergibt, lassen vermuten, dass es sich nur bei der mit A_2^{high} bezeichneten Resonanz um ein 2^+ -Teilchen handelt, während das andere die Quantenzahlen $J^P(A_2^{\text{low}}) = 1^-, 2^+, 3^- \dots$ besitzt. Ein unlängst im CERN ausgeführtes Blasenkammerexperiment²⁷⁾ mit höherer Statistik als das BNL-Experiment findet jedoch in der Reaktion $p\bar{p} \rightarrow K_1 K^\pm \pi^\mp$ bei 0.0, 0.7 und 1.2 GeV Laborenergie des Antiprotons zwei Maxima bei den Massenwerten $M_1 = 1281 \pm 3$ MeV und $M_2 = 1325 \pm 3$ MeV, was für zwei Resonanzen mit den gleichen Quantenzahlen spricht. Aus der empirischen Tatsache, dass für Zweiteilchenreaktionen genügend weit von dem Schwellenwert der totale und unelastische Wirkungsquerschnitt wie $\sigma \sim p_{\text{in}}^{-n}$ vom Eingangsimpuls p_{in} im Laborsystem abhängt und für die Reaktion $\pi^\pm p \rightarrow A_2^\pm p n \approx 0.51 \pm 0.21$ ist, schliesst Morrison²⁸⁾, dass sich an der Stelle des A_2 Mesons mit $J^P = 2^+$ ausserdem ein weiteres Teilchen mit $J^P = 1^+$ oder 2^- befindet. Die Argumente dafür lauten: Empirisch wurde bei anderen Zweiteilchenproduktionsprozessen festgestellt, dass beim Austausch nicht seltsamer Mesonen gilt, dass $n \approx 0$ wenn $\Delta J^{\Delta P}$ (ausgetauschter Drehimpuls und Parität) = $0^+, 1^-, 2^+ \dots$ und $n \approx 1.5$ wenn $\Delta J^{\Delta P} = 0.1, 1^+, 2^- \dots$ ist. Die Erzeugung zweier Teilchen kann somit, als Summe zweier Produktionsprozesse eine Ausnahme zu dieser empirischen Regel vor-täuschen und in einem begrenzten Bereich der Eingangsenergie ein n zwischen 0 und 1.5 ergeben. Gegen dieses Argument spricht der gleiche relative Wirkungsquerschnitt beider Resonanzen bei 2.65 GeV (~ 70 MeV über dem Schwellenwert der Reaktion), 6 und 7 GeV Eingangsenergie²⁹⁾.

Ebenfalls spricht gegen die Hypothese zweier Resonanzen mit verschiedenen Quantenzahlen die Spin-Paritätsanalyse der Daten des Missing-Mass-Experimentes²⁹⁾: Für das gesamte A_2 -Gebiet ($1260 < M < 1360$ MeV) wird $J^P = 2^+$ mit $P(\chi^2) = 40\%$ Wahrscheinlichkeit vor $J^P = 1^-, 1^+$ und 2^- mit

$P(\chi^2) \leq 0.5\%$ eindeutig bevorzugt. Eine getrennte Analyse der beiden Hälften des A_2 - Gebietes ($1254 < M_1 < 1370$ MeV) und ($1207 < M_2 < 1360$ MeV) bevorzugt, ebenfalls - jedoch wegen geringer Statistik nicht sehr signifikant - $J^P = 2^+$ für beide Hälften.

Ein A_2 mit $J^P = 2^+$ und einer Partialbreite $\Gamma(A_2 \rightarrow \rho\pi)$ von 21 ± 5 MeV liesse sich nur schwierig in das SU3 Klassifizierungsschema einordnen: Aus einer Partialbreite für den Zerfall des $K^*(1400)$ in $K^*(890)$ und π sagt die SU3-Theorie nämlich für $A_2 \rightarrow \rho\pi$ eine Zerfallsbreite von 90 MeV voraus.

c) Zwei Resonanzen mit gleichen Quantenzahlen

Die höchste Wahrscheinlichkeit $P(\chi^2) \geq 40\%$ erhält man, wenn man zwei interferierenden Resonanzen an die Daten von Abb. 25 anpasst. Dabei ergeben folgende Parameter die grössten Wahrscheinlichkeiten:

$\alpha)$ Zwei kohärente Breit-Wigner Kurven mit:

$$\begin{aligned} M_1 &= 1289 \text{ MeV}, \quad \Gamma_1 = 22 \text{ MeV} \\ M_2 &= 1309 \text{ MeV}, \quad \Gamma_2 = 22 \text{ MeV} \end{aligned} \quad P(\chi^2) \geq 40\%$$

mit einer relativen Phase von $175 \pm 10^\circ$ und einer gleichen relativen Höhe der Maxima über den Untergrund.

$\beta)$ Der zweiparametrische Ausdruck eines Dipoles³⁰⁾

$$N(M^2) = \frac{\Gamma^2 (M - M_0)^2}{((M - M_0)^2 + \frac{1}{4} \Gamma^2)^2}$$

$$M_0 = 1298 \text{ MeV und } \Gamma = 28 \text{ MeV}$$

gab mit $P(\chi) \cong 50\%$ die grösste Wahrscheinlichkeit.

$\gamma)$ Eine gleich grosse Wahrscheinlichkeit wie bei $\alpha)$ wurde auch durch die kohärente Summe einer breiten und einer schmalen Resonanz bei fast gleicher Masse erreicht. Das Loch im Spektrum entsteht durch die schmale Resonanz, wenn sie destruktiv mit der breiten interferiert. Als beste Parameter ergaben sich:

$$M_1 = 1298 \text{ MeV}, \quad \Gamma_1 = 90 \text{ MeV}$$

$$M_2 = 1297 \text{ MeV}, \quad \Gamma_2 = 12 \text{ MeV}.$$

Bemerkung zu β):

Die Massenverteilung eines Dipols lässt sich aus einem Pol zweiter Ordnung der S-Matrix³⁰⁾ ableiten. Eden und Landshoff³¹⁾ zeigen, dass ein Dipol nur durch das zufällige (und deshalb recht unwahrscheinliche) Zusammenfallen zweier einfacher Resonanzen mit gleichen Quantenzahlen entsteht. Dazu betrachten sie in der Potentialstreuung die S-Matrix als stetige Funktion der Potentialparameter und zeigen, dass dann die Position der Pole stetig von den Potentialparametern abhängen. Unterhalb der Schwelle der Reaktion entsprechen den Polen die gebundenen Zustände: dort sind die Pole streng von erster Ordnung. Die Ordnung der Pole ändert sich nicht, wenn die Potentialparameter so geändert werden, dass die Pole in jenen Bereich der komplexen Ebenen gelangen, wo sie Resonanzen entsprechen. Daraus folgern die Autoren, dass Doppelpole nur durch Zusammenwachsen einzelner Pole entstehen. Diese Lösung unterscheidet sich also nicht grundsätzlich von der in α) oder γ) gefundenen.

Bemerkung zu α) und γ):

Die SU3-Symmetrie verlangt in diesem Fall, dass es an Stelle des einen bisher bekannten 2^+ -Nonets zwei geben sollte. Während das $K^*(1420)$ die gleiche Struktur wie das A_2 zeigen sollte, verlangen die beiden A_2 's zwar jeweils zwei f^0 bzw. f'^0 -Teilchen, die aber nicht in ihrer Masse entartet sein müssen.

D.M. Austin et al.³²⁾ untersuchen im Reggepolmodell die Reaktion $\pi^- p \rightarrow \eta n$, bei der das A_2 -Meson das einzige bekannte Teilchen ist, das ausgetauscht werden kann. Während die Interpretationen, bei denen nur ein einziges 2^+ -Teilchen ausgetauscht wird, eine nichtlineare A_2 -Trajektorie ergeben, kann mit zwei 2^+ -Teilchen eine gute Anpassung an die Daten mit zwei linearen Trajektorien erreicht werden. Eine von beiden (A_2^{high}) fällt im Wesentlichen mit der ρ -Trajektorie zusammen, während die andere (A_2^{low}) den gleichen Schnittpunkt für $t = 0$ wie die π -Trajektorie besitzt.

Herrn Dr. N. Schmitz danke ich herzlich für die Betreuung dieser Arbeit und für wertvolle Anregungen. Herrn Prof. P. Preiswerk möchte ich vielmals für die Möglichkeit danken, diese Arbeit innerhalb der CERN-Bosonspektrometergruppe durchführen zu können.

Mein Dank gilt vor allem auch den Gruppenmitgliedern R. Baud, H. Benz, B. Bošnjaković, D.R. Botterill, G. Damgaard, M.N. Focacci, G. Laverrière, C. Lechanoine, M. Martin, C. Nef, P. Schübelin, A. Weitsch und besonders dem Leiter der Gruppe Herrn Dr. W. Kienzle für die herzliche Aufnahme in der Gruppe.

Für das Tippen der Arbeit und Zeichnen der Abbildungen bin ich Mme Kleinpeter und Mme Lambert zu Dank verpflichtet.

LITERATURANGABEN

- 1) S. Fukui und S. Miyamoto, Nuovo Cimento 11, 113 (1959).
- 2) J. Meek und J.D. Craggs, "Electrical Breakdown of Gases", Oxford Clarendon Press (1953).
- 3) N.A. Kapzow, "Elektrische Vorgänge in Gasen und im Vacuum", VEB-Berlin 1955.
- 4) H. Tholl, Zeitschrift für Naturforschung 18a, S587 (1963).
- 5) G. Charpak, L. Massonet und J. Favier, Progress in Nuclear Techniques and Instrumentation, Band I, S321 (1965), Editor F.J.M. Farley, North-Holland P.C., Amsterdam.
- 6) J.W. Cronin, "Spark Chambers, Bubble and Spark Chambers I, S315 (1967), Editor, R.P. Shutt, Academic Press, London.
- 7) G.E. Chikovani, V.A. Mikhailov, V.N. Roinishvili und A.K. Djavrishvile, Nuclear Instruments and Methods 29, 261 (1964).
B.A. Dolgoshein und B.I. Luchkov, JETP 46, 392 (1964).
- 8) T.H. Romanowski, Diskussion während des Spark-Chamber Symposium, Argonne National Laboratory, Februar 1961, abgedruckt in Rev.Sci.Instr. 32, S516-518 (1961).
- 9) F. Schneider, CERN-Internal Report AR/INT 63-9 (1963).
- 10) V.N. Bolotov und H.I. Devichev, F.I.A.N. Moskau-report A49 (1963).
- 11) G.E. Chikovani, M. Fischer, M.N. Focacci, W. Kienzle, G. Laverrière, C. Lechanoine, B. Levrat, M. Martin und P. Schübelin, CERN 67-13 (1967);
G.E. Chikovani, G. Laverrière und P. Schübelin, Nuclear Instruments und Methods 47, 273 (1967).
- 12) E. Marx, Elektrotechnische Zeitschrift 45, 652 (1924).
E. Gygi und F. Schneider, CERN 64-46 (1964).
- 13) H. Gelerntner, Nuovo Cimento 22, 631 (1961).
- 14) B.C. Maglić und F. Kirsten, Nuclear Instruments and Methods 17, 49 (1962).
- 15) F. Krienen, Nuclear Instruments and Methods 20, 168 (1963).

- 16) V. Perez-Mendez and J. Pfab, UCRL 11, 620 (1964).
- 17) L. Kaufman, UCRL 16, 563 (1966).
- 18) R.M. Bozorth, "Ferromagnetism", van Nostrand Comp. (1951).
- 19) G. Goldhaber, J.L. Brown, S. Goldhaber, J.A. Kadik, B.C. Shen, G.H. Trilling, Phys.Rev.Letters 12, 336 (1964).
- 20) G.E. Chikovani, M.N. Focacci, W. Kienzle, C. Lechanoine, B. Levrat, B. Maglič, M. Martin, P. Schübelin, L. Dubal, M. Fischer, P. Griedar, C. Nef, Phys.Letters 28B, 233 (1968).
- 21) B.C. Maglič und G. Costa, Phys.Letters 18, 185 (1965).
- 22) V. Barger, Rev.Mod.Physics 40, 129 (1968);
J. Stodolsky und J.J. Sakurai, Phys.Rev.Letters 11, 179 (1963).
- 23) H. Benz, G.E. Chikovani, G. Damgaard, M.N. Focacci, W. Kienzle, C. Lechanoine, M. Martin, C. Nef, P. Schübelin, R. Baud, B. Bòšnjaković, J. Cotteron, R. Klanner, A. Weitsch, Phys.Letters 28B, 233 (1968).
- 24) S.U. Chung, O.I. Dahl, L.M. Hardy, R.I. Hess, G.R. Kalbfleisch, J. Kirz, D.H. Miller, G.A. Smith, Phys.Rev.Letters 12, 621 (1964).
- 25) N. Barash-Schmidt, A. Barbaro Galtieri, L.R. Price, A.H. Rosenfeld, P. Söding, Ch.C. Wohl, M. Roos, G. Conforto, Rev.Mod.Physics 40, 77 (1968).
- 26) D.J. Crennell, U. Karshan, K.W. Lai, J.M. Scarr, I.O. Skillicorn, Phys.Rev.Letters 20, 1318 (1968).
- 27) Aguilar Beritez, J. Barlow, L.D. Jacobs, P. Malecki, L. Montand, Ch.d'Andlau, A. Astier, J. Cohen-Ganouna, M. Della Negra, B. Lörstad, "Structure in the $\bar{K}K$ -Decay Mode of A_2 -Meson" in Int.Konferenz für Hochenergiephysik 14, Wien paper 955, (1968).
- 28) D.R.O. Morrison, Phys.Letters 25B, 238 (1967).
- 29) G.E. Chikovani, M.N. Focacci, W. Kienzle, U. Kruse, C. Lechanoine, M. Martin, P. Schübelin, Phys.Letters 28B, 526 (1968).

30) M.L. Goldberger und K.M. Watson, Phys.Rev. 136, B1472 (1964);

M.L. Goldberger und K.M. Watson, Collision Theory, S504,
John Wiley and Sons (1964).

31) Eden und Landshoff, Phys.Rev. 136, B1817 (1964).

32) D.M. Austin, J.V. Beaupre, K.E. Lassila , Phys.Rev. 173, 1573 (1968).

1. The first part of the report is a summary of the work done during the year.

2. The second part is a detailed account of the work done during the year.

3. The third part is a summary of the work done during the year.

4. The fourth part is a summary of the work done during the year.