

.

Einfuehrung in die Theorie dynamischer, insbesondere linearer Systeme und ihre Anwendung in der Psychophysik

U. Mortensen

Fachbereich Psychologie und Sportwissenschaften, Institut III
Westfälische Wilhelms-Universität
Münster, 2000

Der Text soll Hintergrundwissen für mathematische Modellierungen in
der Visuellen Psychophysik liefern.

Revidierte Version 30. 08. 2007
letzte Änderung: 23. 05. 2008, 26.03. 2025

Inhaltsverzeichnis

1	Dynamische Systeme	5
1.1	Dynamische Systeme: allgemeine Begriffe	5
1.2	Lineare Systeme	10
1.2.1	Allgemeine Charakterisierung	10
1.2.2	Die Lösungen eines linearen Systems	13
1.2.3	Die Exponentialmatrix	13
1.2.4	Die Dirac-Impuls-Funktion und die Impulsantwort-Matrix . .	17
1.2.5	Die Operator-Schreibweise	19
1.2.6	Zur Berechnung der Impulsmatrix	22
1.2.7	Stabilität linearer Systeme; Eigenfrequenzen	23
1.2.8	Anmerkungen zum Zustandsbegriff	24
1.2.9	Lineare Systeme n -ter Ordnung	25
1.2.10	System- und Transferfunktionen	34
1.2.11	Minimum-Phasen-Systeme	42
1.2.12	Transferfunktionen und Partialbruchzerlegung	46
1.2.13	Der Fall einfacher Pole	48
1.2.14	Der Fall mehrfacher Pole	49
1.2.15	Elementare Systemelemente	51
1.2.16	Elementare Übertragungsglieder	51
1.2.17	Teilsysteme zweiter Ordnung	54
1.2.18	Graphische Repräsentationen und Maßeinheiten	56
1.2.19	Filter und ihre Charakterisierung	59
1.2.20	Maximale Antwort linearer Filter; Matched Filter	65
1.2.21	Energie und Leistung	67
1.2.22	Linearisierung nichtlinearer Systeme	73

2	Die allgemeine Theorie der Signalentdeckung	82
2.1	Entdeckenswahrscheinlichkeiten als bedingte Wahrscheinlichkeiten . . .	84
2.2	Der Likelihood-Quotient für Weißes Rauschen	86
2.3	Spezielle Modelle des Entdeckens	91
2.3.1	Peak-Detection	91
2.3.2	Wahrscheinlichkeitssummation	92
2.4	Psychometrische Funktionen	93
2.4.1	Allgemeine Modelle	93
2.4.2	Das Quick-Modell: W-Summation zwischen Kanälen	96
2.4.3	Zeitliche Wahrscheinlichkeitssummation	102
3	Visuelle Psychophysik	109
3.1	Psychophysik im Zeit- und Ortsbereich	109
3.2	Erste empirische Befunde	110
3.3	Systemidentifikation im Zeitbereich	111
3.3.1	Schätzung der Amplitudencharakteristik	115
3.3.2	Die Perturbationsmethode	122
3.3.3	Peak-Detection versus Wahrscheinlichkeitssummation	125
3.4	Systemidentifikation im Ortsbereich	130
3.4.1	Das visuelle System als einzelner Kanal	130
3.4.2	Die direkte Schätzung der Impulsantwort im Ortsbereich . . .	133
3.4.3	Neuronale Kanäle: Orientierungen	136
3.4.4	Neuronale Kanäle: Ortsfrequenzen	139
3.4.5	Neuronale Kanäle: Kanten-, Linien- und Gitterdetektoren . .	141
3.4.6	Neuronale Kanäle als Matched Filter	146
A	Anhang: Fourier- und Laplace-Transformationen	154
A.1	Komplexe Zahlen	154
A.2	Fourier-Reihen	158
A.2.1	Periodische Funktionen	158
A.2.2	Superpositionen	159
A.2.3	Die Orthogonalitätsrelationen der Sinus- und Kosinusfunktionen	162
A.3	Fouriertransformierte	164
A.3.1	Definition der Fouriertransformierten	164
A.3.2	Einige allgemeine Sätze	166

A.3.3	Die Dirac-Delta Funktion	171
A.4	Fourier-Transformationen spezieller Funktionen	172
A.4.1	Die Sinus- bzw. Kosinusfunktion	172
A.4.2	Die Exponentialfunktion	172
A.4.3	Eine Konstante	173
A.4.4	Die Funktion $1/x$	175
A.4.5	Die Schrittfunktion	175
A.4.6	Der Balken	176
A.4.7	Die Gabor-Funktion	177
A.4.8	Hermite-Funktionen	178
A.5	Hilbertransformierte	179
A.6	Laplace-Transformierte	181
A.7	Die Schwartzsche Ungleichung	182
A.8	2-dimensionale Signale	183
A.8.1	Allgemeine Definition	183
A.8.2	Gabor-Muster	185
A.9	Stochastische Prozesse	185
A.10	Die Weibull-Verteilung	194
Literatur		195
Index		201

Symbole und Bezeichnungen

Es wird von den folgenden Symbolen bzw. Bezeichnungen Gebrauch gemacht:

$\mathbb{R}$	Menge der reellen Zahlen
$\mathbb{N}$	Menge der natürlichen Zahlen $0, 1, 2, \dots$
$\mathbb{C}$	Menge der komplexen Zahlen $z = a + ib$, $a, b \in \mathbb{R}$, $i = \sqrt{-1}$
$\exp(x)$	e^x
$\log$	natürlicher Logarithmus (zur Basis e)
$\log_{10}$	Logarithmus zur Basis 10

Kapitel 1

Dynamische Systeme

1.1 Dynamische Systeme: allgemeine Begriffe

Sprungfedern, Planetensysteme, Pendel, interagierende Populationen wie die der Füchse und Hasen, Neurone und Netze von Neuronen sind Beispiele für dynamische Systeme; diese Systeme müssen also trotz ihrer Verschiedenheit wesentliche Aspekte gemein haben. Die Definition des Begriffs des dynamischen Systems wird dementsprechend einigermaßen abstrakt sein. Es soll gleichwohl auf allzu große mathematische Strenge verzichtet werden, da sie den Rahmen des Textes sprengen würde; der interessierte Leser sei auf die Literatur, etwa Arnol'd (1980), Jänich (1983) oder Arrowsmith und Place (1990) verwiesen. Es genügt für unsere Zwecke, ein gewisses intuitives Verständnis der Begriffsbildungen zu erhalten, das die qualitative Betrachtung dynamischer Systeme erleichtern soll.

Ohne Beschränkung der Allgemeinheit läßt sich sagen, dass ein System zu einem bestimmten Zeitpunkt in einem bestimmten Zustand ist; so hat ein schwingendes Pendel zu einem Zeitpunkt t eine bestimmte Auslenkung $\phi(t)$ (Abweichung von der Vertikalen) sowie eine bestimmte Geschwindigkeit $v(t)$, und $\phi(t)$ und $v(t)$ zusammen definieren den Zustand des Pendels. Man sieht hier bereits, dass bei der Charakterisierung der Zustände eine gewisse Willkürlichkeit herrscht, denn von der Beschaffenheit des Materials, von seinem Alter und einer Reihe anderer Eigenschaften sehen wir bei unserer Zustandsbeschreibung ab. Betrachten man ein Neuron, so kann sein Membranpotential zur Zeit t als sein Zustand zur Zeit t erklärt werden, und wieder hat man dabei von vielen Eigenschaften des Neurons, die man ebenso als zustandsdefinierend betrachten könnte, abstrahiert. Die Wahl der betrachteten Zustände ist somit ein Teil der Theorie, die über das System entwickelt wird. Für die Definition des Begriffs des dynamischen Systems ist dieser Aspekt der Willkürlichkeit belanglos.

Der Zustand des betrachteten Systems im Zeitpunkt t soll durch $x(t)$ repräsentiert werden, wobei $x(t)$ als Vektor

$$x(t) = (x_1(t), \dots, x_n(t))'$$

aufgefaßt wird. Die n Komponenten $x_i(t)$, $1 \leq i \leq n$, von $x(t)$ repräsentieren dann

bestimmte Zustandsaspekte und können als Koordinaten eines Punktes x bzw. $x(t)$ im $\mathbb{R}^n$ aufgefaßt werden. Dementsprechend werden wir die Ausdrücke "Punkt" und "Zustand" synonym verwenden. Die $x_i(t)$ sind Funktionen der Zeit. Der Verlauf dieser Zustandsänderungen definiert die Dynamik des Systems. Die Definition des Begriffs des dynamischen Systems wird dementsprechend auf einen "Zustandsraum" als Teilraum des $\mathbb{R}^n$ bezug nehmen und eine allgemeine Charakterisierung der Dynamik enthalten.

Ist nun ein Zustand durch einen Punkt $x(t)$ gegeben, so bedeutet die Veränderung des Zustands nach s Zeiteinheiten den Übergang zu einem anderen Punkt $x(t+s) = (x_1(t+s), \dots, x_n(t+s))'$. Formal wird also die Dynamik durch eine Abbildung von Punkten im $\mathbb{R}^n$ auf Punkte im $\mathbb{R}^n$ erklärt. Dementsprechend hat man dann die

Definition 1.1.1 *Es sei $M \subset \mathbb{R}^n$ ein offenes Gebiet¹ des $\mathbb{R}^n$ und $\phi: \mathbb{R} \times M \rightarrow M$ sei eine stetige Abbildung, die für alle $x \in M$ den Bedingungen*

$$(i) \quad \phi(0, x) = x$$

$$(ii) \quad \phi(s, \phi(t, x)) = \phi(t+s, x) \text{ für alle } s, t \in \mathbb{R}$$

genügt. Dann heißt ϕ dynamisches System, und c heißt Phasen- oder Zustandsraum. $x \in M$ heißt Zustand oder Phasenpunkt. Das Paar (M, ϕ) heißt Zustands- oder Phasenfluß.

Gemäß Definition 1.1.1 ist ein dynamisches System abstrakt als eine Abbildung ϕ definiert, und zwar werden Elemente von $\mathbb{R} \times M$ auf Punkte in c abgebildet. Sind c und ϕ gegeben, so hat man das System spezifiziert. Ein Element aus $\mathbb{R} \times M$ ist ein Paar (t, x) , durch das ein Zeitpunkt und ein Zustand festgelegt werden. Die Abbildung ϕ ordnet diesem Paar, d.h. dem Zustand x zum Zeitpunkt t , einen bestimmten anderen Zustand $y \in M$ zu, nämlich $y = \phi(t, x)$. Durch ϕ werden die Zustandstransformationen oder der Fluß der Zustände festgelegt; daher der Ausdruck Zustandsfluß für das Paar (M, ϕ) . (Der synonym gebrauchte Ausdruck *Phase* für Zustand geht auf Betrachtungen in der Physik zurück.) Die Bedingungen (i) und (ii) in Definition 1.1.1 heißen dementsprechend auch *Flußaxiome*.

Für die in Definition 1.1.1 eingeführte Abbildung ϕ ist Stetigkeit gefordert worden; dies bedeutet, dass kein sprunghafter Übergang von einem Zustand zu einem anderen erfolgen soll; jeder Übergang benötigt Zeit. Die Implikationen dieser Forderung werden anhand des in der folgenden Definition eingeführten Begriffs, der es erlaubt, weitere Eigenschaften eines dynamischen Systems herzuleiten, deutlich werden.

Definition 1.1.2 *Es sei $M \subset \mathbb{R}^n$ der Zustandsraum eines dynamischen Systems und es sei $\varphi_x: \mathbb{R} \rightarrow M$, $\varphi_x(t) \mapsto y \in M$, eine Abbildung, die für spezielles $x \in M$ und $t \in \mathbb{R}$ den Zustand $y \in M$ festlegt, den das System von x ausgehend nach t Zeiteinheiten annimmt. Dann heißt φ_x Flußlinie oder Bewegung des Punktes x unter*

¹Die Randpunkte des Gebietes gehören nicht zu M .

Wirkung des Flusses (M, ϕ) . Die Menge $\varphi_x(\mathbb{R})$ der Punkte, die man von x ausgehend erreicht, heißt die durch x gehende Trajektorie (Bahn, Orbit, Phasenkurve) des dynamischen Systems ϕ .

Eine Trajektorie des dynamischen Systems ϕ ist dementsprechend das *Bild* der Abbildung φ_x . Natürlich ist $\varphi_x(t) = \phi(t, x)$. Während aber ϕ für *beliebiges* x und t den Zustand y definiert, in den x nach t Zeiteinheiten übergeht, betrachtet φ_x diesen Übergang nur für ein *spezielles* $x \in M$, d.h. betrachtet man φ_x , so betrachtet man einen speziellen, durch den Punkt x verlaufenden Prozeß. Dementsprechend ist die Kenntnis *aller* Abbildungen, d.h. aller Flußlinien φ_x gleichbedeutend mit der Kenntnis von ϕ . Man erhält nun sofort die im folgenden Satz enthaltenen Folgerungen:

Satz 1.1.1 *Es gelten die folgenden Aussagen:*

1. $\varphi_x(0) = x$ für alle $x \in M$, d.h. der Zustand verändert sich nicht, ohne dass Zeit vergeht.
2. Es sei $y = \varphi_x(t_0) \in M$; dann läßt sich auch eine Flußlinie für y erklären. Offenbar ist $\varphi_y(t) = \varphi_x(t_0 + t)$ für alle $t \in \mathbb{R}$.
3. Haben zwei Flußlinien φ_x und φ_y einen gemeinsamen Punkt, so sind sie überhaupt identisch; zwei nicht identische Flußlinien können sich nicht kreuzen.

Die Folgerungen 1. und 2. sind unmittelbar klar. Die Folgerung 3. sieht man wie folgt ein: es gelte etwa $\varphi_y(t) = \varphi_x(t_0 + t)$. Dies heißt aber $y = \varphi_x(t_0)$, und nach Satz 1.1.1 erreicht man alle Zustände $z = \varphi_y(t)$ ebenfalls, wenn man von x ausgehend die Zustände $\varphi_x(t_0 + t)$ bestimmt. Dies heißt, dass die Menge $\varphi_y(\mathbb{R})$ der Punkte, die man von y ausgehend erreicht, gleich der Menge $\varphi_x(\mathbb{R})$ der Punkte ist, die man von x ausgehend erreicht; die durch x und y gehenden Trajektorien sind also gleich. Da sich aus $\varphi_y(t) = \varphi_x(t_0 + t)$ die Gleichheit der Trajektorien ergibt, folgt aus der *Ungleichheit* der Trajektorien $\varphi_x(\mathbb{R})$ und $\varphi_y(\mathbb{R})$ auch, dass es keine Punkte x und y geben kann, für die $\varphi_y(t) = \varphi_x(t_0 + t)$ gilt. Durch jeden Punkt $x \in M$ kann also *nur* eine Trajektorie gehen.

Es sollen nun einige Typen von Trajektorien charakterisiert werden; anschließend wird gezeigt, dass es nur diese Typen gibt:

Definition 1.1.3 *Gilt $\varphi_x(t) = x$ für alle $t \in \mathbb{R}$, so heißt x ein Fixpunkt des durch ϕ definierten dynamischen Systems. Existiert ein $\tau_0 = \inf\{t_0 | \varphi_x(t + t_0) = \varphi_x(t_0)\}$ ², $\tau_0 > 0$, derart, dass $\varphi_x(t + \tau_0) = \varphi_x(t)$, so heißt φ_x periodisch und τ_0 dementsprechend die Periode. Gilt $\varphi_x(t_0) \neq \varphi_x(t_1)$ für $t_0 \neq t_1$, so heißt φ_x injektiv.³*

Die Bedeutung dieser Begriffe wird in der Aussage des folgenden Satzes verdeutlicht, die auch für qualitative Betrachtungen dynamischer Systeme von Bedeutung ist:

²*inf* steht für *infimum*; dies ist in diesem Fall die kleinste Zahl der in $\{\dots\}$ angegebenen Menge. Dementsprechend soll also $\tau_0 \leq t_0$ für alle t_0 gelten.

³Allgemein heißt eine Abbildung $f : A \rightarrow B$ *injektiv*, wenn für alle x_1 und x_2 aus A folgt, dass $x_1 = x_2$ falls $f(x_1) = f(x_2)$ ist. In Definition 1.1.2 ist aber φ_x als Abbildung definiert worden.

Satz 1.1.2 *Die Flußlinien des Systems sind entweder Fixpunkte, oder sie sind periodisch, oder sie sind injektiv.*

Beweis: Den Beweis des Satzes findet man etwa in Arnol'd (1980), Jänich (1983), Arrowsmith und Place (1989).

In der Abb. 1 sind die drei Typen von Flußlinien dargestellt; Abb. 2 enthält die Flußlinien, die sich bei einem ebenen Pendel ergeben. Die wie eine Acht aussehende Trajektorie darf nicht als sich kreuzende Flußlinie aufgefaßt werden, denn nach Satz 1.1.2 gibt es keine Flußlinien, die sich kreuzen; der "Kreuzungspunkt" ist in der Tat eine Flußlinie für sich, nämlich ein *Fixpunkt*; gerät das System in diesen Zustand, so verharrt es dort, d.h. es finden keine Veränderungen der x_i statt. Beim Pendel ist dieser Zustand z.B. dann erreicht, wenn es zur Ruhe gekommen ist und senkrecht herunterhängt. Der Fixpunkt definiert einen *Gleichgewichtszustand*.

Die Definition des dynamischen Systems durch die Abbildung ϕ und die der Flußlinien φ_x implizieren also, dass die Bewegungen im Zustands- oder Phasenraum in nur drei Klassen zerlegt werden können. Die Tatsache, dass es für einen Punkt (Zustand) $x \in M$ nur eine Trajektorie geben kann, die durch ihn hindurchführt, bedeutet, dass die Kenntnis von x sofort auch die Kenntnis des Prozesses, d.h. des Verlaufs der Übergänge der Zustände ineinander, bedeutet. Der Prozeß ist insgesamt eindeutig bestimmt oder *determiniert*.

Bis jetzt ist zwar ein dynamisches System als eine Abbildung ϕ definiert worden, aber es ist noch unklar, wie man zu einer konkreten Charakterisierung des Verhaltens eines solchen Systems gelangen kann. Mit *Verhalten* sind dabei die Bewegungen im Zustandsraum gemeint. Nun bedeutet die Kenntnis von ϕ nach dem bisher Gesagten auch eine Kenntnis des Verhaltens, und so kann man sagen, dass ϕ eine funktionale Beschreibung des Systems ist. Umgekehrt ist man häufig daran interessiert, aus Beobachtungen des Verhaltens auf eine funktionale, durch Definition einer Abbildung ϕ gegebene Beschreibung des Systems zu kommen. Möglichkeiten, ϕ näher zu charakterisieren, sollen jetzt angegeben werden.

Die Dynamik eines Systems äußert sich in den Bewegungen der Zustandspunkte im Phasenraum M . Die Bewegungen, d.h. die Zustandsveränderungen, können, abhängig vom Ort in c schnell oder langsam sein. Wir können also jedem Punkt $x \in M$ eine Geschwindigkeit $v(x)$ zuordnen. Denn durch x geht nur eine Flußlinie, und deshalb kann es in x auch nur eine Geschwindigkeit $v(x)$ geben. $v(x)$ ist ein Maß für die Veränderung der Flußlinie in x und es liegt deshalb nahe, sie über einen Differentialquotienten zu definieren:

Definition 1.1.4 *Die Ableitung*

$$v(x) = d\varphi_x(t)/dt|_{t=0}$$

heißt Phasengeschwindigkeit des Flusses (M, ϕ) im Punkt x . Die Abbildung $v : M \rightarrow \mathbb{R}^n$ heißt Vektorfeld auf M .

Der Punkt x ist durch einen Vektor definiert, und deshalb ist $v(x)$ ebenfalls ein Vektor. Dies kommt noch einmal in der Definition des Vektorfeldes v zum Ausdruck. Jedem Punkt in M wird ein Punkt in $\mathbb{R}^n$ zugeordnet, dessen Koordinaten eben die Komponenten des Vektors v sind.

$x \in M$ ist ein Zustand des Systems, und φ_x beschreibt die Transformation von x als Funktion von t . $v(x)$ gibt dann die Veränderung der Flußlinie oder Trajektorie in x an. Dementsprechend kann man auch $\dot{x}$ für $v(x)$ schreiben. Man hat dann die

Definition 1.1.5 *Es sei $M \subset \mathbb{R}^n$ ein Zustandsraum eines dynamischen Systems ϕ und v das ϕ entsprechende Vektorfeld. Dann heißt*

$$\dot{x} = v(x) \tag{1.1}$$

die durch das Vektorfeld v definierte Differentialgleichung (abgekürzt: Dgl). c heißt der Phasen- oder Zustandsraum der Dgl (1.1). Die Flußlinie φ_x heißt Lösung der Dgl (1.1), wenn $d\varphi_x(t)/dt|_{t=a} = v(\varphi_x(\tau))$ für $\tau = t$, $a < t < b$, wobei $a = -\infty$ und $b = +\infty$ zugelassen sind. Gilt $\varphi_x(t_0) = x_0 \in M$ für $a < t_0 < b$, so sagt man, dass die Lösung der Anfangsbedingung $\varphi_x(t_0) = x_0$ genügt.

Die Dgl (1.1) ist ein System von Differentialgleichungen; ausgeschrieben erhält man

$$\begin{aligned} \dot{x}_1 &= v_1(x_1, \dots, x_n) \\ \dot{x}_2 &= v_2(x_1, \dots, x_n) \\ &\vdots \\ \dot{x}_n &= v_n(x_1, \dots, x_n) \end{aligned} \tag{1.2}$$

Die v_i sind dabei Funktionen der $x_j = x_j(t)$. Ist das System ϕ bekannt, so kann es durch das zugehörige System von Dgln (1.2) beschrieben werden. Auf der linken Seite treten nur Ableitungen erster Ordnung auf, und auf der rechten Seite findet man nur Funktionen der x_i , nicht aber ihre Ableitungen:

Ein dynamisches System ist also stets durch ein System von Differentialgleichungen erster Ordnung repräsentierbar.

Es sollte an hier Stelle angemerkt werden, dass durch (1.1) *deterministische* dynamische Systeme beschrieben werden. Stochastische Irregularitäten des Systemverhaltens gehen bei der Beschreibung durch (1.1) nicht mit ein. Das Verhalten neuronaler Systeme enthält aber stochastische Komponenten, so dass eigentlich stochastische dynamische Systeme betrachtet werden müssen. Die Diskussion solcher Systeme setzt aber die Kenntnis des Begriffs des deterministischen Systems voraus. Auf die stochastischen Aspekte wird später eingegangen, es genügt an dieser Stelle, sich auf die Klasse der deterministischen Systeme zu beschränken.

Ist ein dynamisches System durch ein System von Dgln charakterisierbar, so kann man umgekehrt *Annahmen* über ein System durch die Formulierung entsprechender Dgln artikulieren. Bestimmte — nicht alle — Systeme lassen sich auch durch bestimmte Funktionen, die Impuls- und die Transferfunktion, charakterisieren; darauf

wird später ausführlich eingegangen, denn von dieser Möglichkeit wird bei Anwendungen der Theorie der Systeme auf die Psychophysik ausgiebig Gebrauch gemacht. Impuls- und Transferfunktion stehen allerdings in einer bestimmten Beziehung zu dem System von Dgln, durch die das System beschrieben werden soll. Damit erläutert werden kann, wie die empirisch bestimmten Impuls- und Transferfunktionen interpretiert werden können, sollen zunächst die Systeme von Dgln weiter betrachtet werden.

Eine *Lösung* für das System von Dgln besteht in einem Vektor

$$x(t) = (x_1(t), \dots, x_n(t))',$$

dessen Komponenten $x_j(t)$ Funktionen sind, die den in Bedingung (1.2) genannten Forderungen genügen. Die $x_1(t), \dots, x_n(t)$ können als Koordinaten eines Punktes aufgefaßt werden, der sich in Abhängigkeit von der Zeit durch den Phasenraum bewegt. Die Bahn eines solchen Punktes ist eine Trajektorie, für die die Aussage des Satzes 1.2 gilt. Nach Satz 1.1.1, Aussage 3 können sich zwei nicht identische Flußlinien nicht kreuzen; durch jeden Punkt des Phasenraumes kann nur eine Trajektorie gehen. Hat man für die Lösung $x(t)$ also einen Punkt — die *Anfangsbedingung* — spezifiziert, so ist damit der Verlauf von $x(t)$ festgelegt.

Die Funktionen v_i in (1.1) können *linear* oder *nichtlinear* sein. Sind alle Funktionen $v_1, \dots, v_n$ linear, so heißt das System *linear*. Lineare Systeme spielen in den Anwendungen der Systemtheorie eine zentrale Rolle, u.a. deswegen, weil sich die Lösungen für solche Systeme zumindest grundsätzlich anschreiben lassen. Bei nichtlinearen Systemen läßt sich in einfachen Fällen $x(t)$ explizit herleiten, in vielen Fällen ist aber eine solche explizite Herleitung nicht nur schwierig, sondern unnötig: gemäß einem Satz von Liouville existieren vielfach gar keine geschlossenen Ausdrücke für $x(t)$, und wenn sie existieren, sind sie nicht notwendig auch hilfreich, um das Verhalten des Systems zu erfassen (Arnol'd (1980), p. 39). Nun sind die meisten biologischen Systeme, insbesondere auch das visuelle System, nichtlinear. Gleichwohl ist der eben genannte Sachverhalt, dass sich nämlich nur selten die Lösungen in expliziter Form angeben lassen, kein Grund zur Melancholie. Denn für eine große Klasse nichtlinearer Systeme lassen sich *Linearisierungen* angeben, die das Verhalten des Systems für kleine Auslenkungen, d.h. für Störungen mit geringer Intensität, im Rahmen der Theorie linearer Systeme zu beschreiben gestatten. Zumindest im Prinzip ist es auch möglich, von den linearisierten Versionen auf die Nichtlinearitäten zu schließen. Im Folgenden sollen deshalb die linearen Systeme ausführlicher dargestellt werden; Bedingungen für die Linearisierbarkeit eines nichtlinearen Systems werden im Anschluß daran angegeben.

1.2 Lineare Systeme

1.2.1 Allgemeine Charakterisierung

Es sei $\dot{x} = v(x)$ ein System von Dgln. Dann sind x und $\dot{x}$ n -dimensionale Vektoren, und v kann abstrakt als Abbildung $v : \mathbb{R}^n \rightarrow \mathbb{R}^n$ verstanden werden. Ist v

linear, so heißt $\dot{x} = v(x)$ ein *lineares System* mit der Dimension n . Dann kann das System in der Form

$$\dot{x} = v(x) = Ax + u \quad (1.3)$$

geschrieben werden⁴ Dabei ist

$$x = (x_1, \dots, x_n)'$$

der Zustandsvektor, die Komponenten $x_i = x_i(t)$ sind natürlich Funktionen der Zeit und heißen *Zustandsvariablen*. Diese Variablen können direkt physikalische oder psychologische Variable abbilden, sie können aber auch bestimmte Funktionen solcher Variablen sein; auf spezielle Interpretationen wird später eingegangen. A ist die *Koeffizientenmatrix*

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ & & \vdots & \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{pmatrix}$$

Die a_{ij} können Funktionen der Zeit sein, sie können aber auch Konstante sein; im letzteren Fall handelt es sich bei (1.3) um ein *System mit konstanten Koeffizienten*; die folgenden Betrachtungen gelten für beide Arten von Systemen. Der Vektor $u = (u_1(t), \dots, u_n(t))'$ repräsentiert "Störungen" des Systems. Für die i -te Komponente des Vektors x gilt dann jedenfalls

$$\dot{x}_i(t) = \frac{dX_i(t)}{dt} = \sum_{j=1}^n a_{ij}x_j(t), \quad i = 1, \dots, n. \quad (1.4)$$

Für $u(t) \equiv 0$ im betrachteten Zeitintervall ergibt sich

$$\dot{x} = Ax; \quad (1.5)$$

diese Gleichung beschreibt ein *homogenes System*. Ein homogenes System bewegt sich nicht unter dem Einfluß einer äußeren "Störung". Für den Spezialfall $\dot{x} = 0$ findet keine Zustandsänderung statt, d.h. das System führt keinerlei Bewegung aus. Das System ist dann in einer *Gleichgewichtslage* oder in einem *Fixpunkt*. Die Menge der Fixpunkte ergibt sich aus der Diskussion der Gleichung

$$Ax = 0 \quad (1.6)$$

Ist der Rang von A gleich n , d.h. ist $\det A = |A| \neq 0$, so weiß man aus der Theorie der linearen Gleichungssysteme, dass es für (1.6) nur eine Lösung gibt, nämlich $x = 0$ für alle t . Ist dagegen $|A| = 0$, d.h. ist der Rang von A kleiner als n , so sind die Gleichungen (1.6) voneinander abhängig und es existiert eine Hyperebene, auf der (1.6) erfüllt ist. Eine ausführliche Diskussion dieses Falles kann hier ni Knobloch & Kappel (1974)). Es genügt hier, sich auf den Fall $|A| \neq 0$ zu beschränken. Die Gleichungen in (1.6) sind dann voneinander unabhängig und stellen insofern eine

⁴Die Bezeichnung *lineares System* ist üblich, wenn auch ein wenig lax, denn linear (im mathematischen Sinne) ist es nur für den Fall $u \equiv 0$. (1.3) heißt eigentlich *affines System*.

ökonomische Beschreibung des Systems dar; da das Ziel psychophysischer Untersuchungen eine Beschreibung des visuellen Systems oder Teilen davon ist, wird man immer auf eine ökonomische Beschreibung abzielen.

Es sei also $|A| \neq 0$ und es sei t_0 der Anfangspunkt des Intervalles, in dem $u(t) \equiv 0$ gilt, und es sei $x_0 = x(t_0)$. Ist $x_0 \neq 0$, so ist das System *ausgelenkt*, d.h. es ist durch eine vorangehende Störung aus der Gleichgewichtslage in den Zustand x_0 gebracht worden. Man hat etwa durch einen Stoß die Stimmgabel in Schwingungen versetzt, oder die Spiralfeder zu einem gewissen Grade zusammengedrückt oder auseinander gezogen. Die Stimmgabel wird irgendwann schwingen, sie strebt dem Fixpunkt $x = 0$ zu. Die innere Reibung in der Gabel bedeutet, dass die Schwingungen gedämpft werden, bis eben keine Schwingung mehr vorhanden ist. Für ein beliebiges System kann es sein, dass es, von t_0 an sich selbst überlassen, nicht gegen einen Fixpunkt strebt: einige oder alle Komponenten von x können gegen Unendlich streben, oder das System schwingt ohne jede Dämpfung. Strebt x gegen einen Fixpunkt, so ist das System *stabil* *nstabil*. Es zeigt sich, dass es von den Eigenwerten von A abhängt, ob ein System stabil ist oder nicht; hierauf wird in Abschnitt 1.2.7 näher eingegangen.

Bei den in diesem Buch behandelten Systemen handelt es sich um Informationsübertragungssysteme. Dementsprechend kann der Vektor u auch als Vektor von *Eingangssignalen* oder als *Input-Vektor* aufgefaßt werden. Unter bestimmten Bedingungen sind dann die Komponenten des Vektors x die Ausgangssignale. Es wird sich aber oft als nützlich erweisen, die Ausgangssignale als eine bestimmte Transformation des Vektors aufzufassen. So sei C eine $m \times n$ -Matrix und $y = (y_1, \dots, y_m)'$ sei ein c -dimensionaler Vektor, der die Ausgangssignale repräsentieren möge. Außerdem kann man noch zulassen, dass der Eingangsvektor u eine direkte Wirkung auf y haben darf. Dazu sei D eine $m \times n$ -Matrix, und der Ausgangsvektor sei durch $Cx + Du$ definiert. Außerdem können wir noch zulassen, dass der Vektor u in transformierter Form auf den Zustand des Systems einwirkt, d.h. in der Form Bu . Das System (1.3) kann dann in der Form

$$\begin{aligned}\dot{x} &= Ax + Bu \\ y &= Cx + Du\end{aligned}\tag{1.7}$$

angeschrieben werden. Die Matrix B heißt auch *Kontrollmatrix*, C heißt *Output-Matrix*, und D heißt *Feedforward-Matrix*. y heißt *Output-Vektor*.

Das visuelle System ist wie die meisten biologischen Systeme nichtlinear. Die Bedeutung der linearen Systeme ergibt sich aus dem Umstand, dass einerseits lineare Systeme einer mathematischen Analyse leichter zugänglich sind und nichtlineare Systeme für kleine Auslenkungen, d.h. Störungen, häufig durch lineare Systeme in sehr guter Näherung beschrieben werden können. Die Linearisierung nichtlinearer Systeme setzt die Kenntnis einer Reihe von Eigenschaften linearer Systeme voraus, weshalb lineare Systeme in einiger Ausführlichkeit behandelt werden sollen.

Die Gleichung (1.3) ergibt sich direkt aus der Definition des Begriffes Dynamisches System. Dies bedeutet, dass sich *jedes* lineare System in der Form (1.3) darstellen läßt. Deshalb heißt das System von Dgln (1.3) bzw. (1.7) auch die *Normalform*.

1.2.2 Die Lösungen eines linearen Systems

1.2.3 Die Exponentialmatrix

Zur Vorbereitung wird der Begriff der *Exponentialmatrix* eingeführt. Es sei zunächst daran erinnert, dass eine hinreichend oft differenzierbare Funktion $f(x)$ in eine Taylor-Reihe entwickelt werden kann

$$f(x) = f(0) + \frac{x}{1!}f'(0) + \frac{x^2}{2!}f''(0) + \frac{x^3}{3!}f'''(0) + \dots \quad (1.8)$$

Dabei ist

$$f^{(n)}(x) = \frac{d^n}{dx} f(x), \quad n = 1, 2, \dots$$

die n -te Ableitung von f an der Stelle x ; $f^{(n)}(0)$ ist also die n -te Ableitung von f an der Stelle $x = 0$. Setzt man $\xi = x + h$, so kann man f in eine Taylor-Reihe um den Punkt x entwickeln:

$$f(\xi) = f(x + h) = f(x) + \frac{h}{1!}f'(x) + \frac{h^2}{2!}f''(x) + \frac{h^3}{3!}f'''(x) + \dots \quad (1.9)$$

Für hinreichend kleine Werte von h kann also f durch eine *lineare Funktion*

$$f(x + h) \approx f(x) + hf'(x), \quad f'(x) = \frac{df(x)}{dx}$$

angenähert werden; es ist ja $h/1! = h$. $f'(x)$ ist die Steigung, $f(x)$ ist die additive Konstante bei dieser Approximation. Auf diese Approximation wird bei der Linearisierung nichtlinearer Systeme zurückgegriffen.

Es sei P eine $n \times n$ -Matrix. Dann ist $P^2 = P \cdot P$, $P^3 = P \cdot P \cdot P$, etc. Weiter sei a ein Skalar, d.h. eine reelle Zahl ($a \in \mathbb{R}$). Dann ist aP eine Matrix, deren Elemente sich von den Elementen der Matrix P nur durch den Faktor a unterscheiden. Ist also p_{ij} das Element in der i -ten Zeile und j -ten Spalte von P , so ist ap_{ij} das Element in der i -ten Zeile und j -ten Spalte von aP . Man kann dann die unendliche Reihe

$$\frac{P^0}{0!} + \frac{P^1}{1!} + \frac{P^2}{2!} + \dots = \sum_{k=0}^{\infty} \frac{P^k}{k!} \quad (1.10)$$

bilden. Dabei ist P^0 eine Matrix, deren Elemente alle gleich 1 sind, und da $0! = 1$, ist $P^0/0!$ eine Matrix, deren Elemente alle gleich 1 sind. Der Faktor $1/k!$ ist eine reelle Zahl, so dass $P^k/k!$ eine Matrix ist, deren Elemente dadurch entstehen, dass zunächst P k -mal mit sich selbst multipliziert wird und dann die Elemente von P^k mit $a = 1/k!$ multipliziert werden.

Die Reihe (1.10) entspricht *formal* der Taylor-Reihe für die Funktion $f(x) = e^x$, d.h.

$$e^x = 1 + \frac{x^1}{1!} + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots = \sum_{k=0}^{\infty} \frac{x^k}{k!}$$

Wegen der Analogie zu der in (1.10) eingeführten Reihe wird man zu der folgenden Definition geführt:

Definition 1.2.1 *Die Matrix*

$$H = \sum_{k=0}^{\infty} \frac{P^k}{k!} =: e^P \quad (1.11)$$

heißt Exponentialmatrix⁵.

Ein im folgenden Abschnitt interessierender Spezialfall ist die Exponentialmatrix der Matrix At ; es gilt

$$e^{At} = \sum_{k=0}^{\infty} \frac{(At)^k}{k!} \quad (1.12)$$

und

$$e^{-At} = \sum_{k=0}^{\infty} \frac{(-At)^k}{k!} = \sum_{k=0}^{\infty} (-1)^k \frac{(At)^k}{k!} \quad (1.13)$$

Für spätere Anwendungen sei gleich noch der Ausdruck für die Ableitung von e^A bzw. e^{-At} nach t angefügt: es ist

$$\begin{aligned} \frac{d}{dt} e^{At} &= \frac{d}{dt} \left(1 + At + \frac{(At)^2}{2!} + \frac{(At)^3}{3!} + \dots \right) \\ &= A + \frac{2(At)A}{2!} + \frac{3(At)^2 A}{3!} + \frac{4(At)^3 A}{4!} + \dots \\ &= \left(1 + At + \frac{(At)^2}{2!} + \frac{(At)^3}{3!} + \dots \right) A \\ &= e^{At} A \end{aligned} \quad (1.14)$$

Andererseits gilt für die Terme $(At)^{k-1}A/(k-1)!$

$$\frac{k(At)^{k-1}A}{k!} = \frac{(At)^{k-1}A}{(k-1)!} = \frac{A^k t^{k-1}}{(k-1)!} = A \frac{(At)^{k-1}}{(k-1)!},$$

denn t ist ja ein Skalar, d.h. es gilt $tA = At$, so dass man auch

$$\frac{d}{dt} e^{At} = A e^{At} \quad (1.15)$$

schreiben kann. Auf analoge Weise zeigt man, dass

$$\frac{d}{dt} e^{-At} = -A e^{-At} \quad (1.16)$$

ist.

⁵Das Zeichen "=: " soll heißen, dass das Symbol e^P durch die links stehende Summe definiert wird.

Herleitung der Lösung

Es wird zunächst das homogene System $\dot{x} = Ax$ betrachtet. Es gilt der

Satz 1.2.1 *Es sei A eine konstante Matrix, d.h. die Elemente a_{ij} der Matrix A sollen nicht von der Zeit abhängen. Dann ist $x(t)$ durch*

$$x(t) = x_0 e^{At} \quad (1.17)$$

gegeben.

Beweis: Eine Möglichkeit, die Lösung für das homogene System herzuleiten, besteht darin, den Lösungsvektor $x(t)$ in eine Taylor- bzw. MacLaurin-Reihe zu entwickeln; dies bedeutet, dass jede Komponente von x in eine solche Reihe entwickelt wird. Mit $x_0 = x(0)$ soll also gelten

$$\begin{aligned} x(t) &= x_0 + t x^{(1)}(0) + \\ &\quad + \frac{t^2}{2!} x^{(2)}(0) + \frac{t^3}{3!} x^{(3)}(0) + \dots \end{aligned} \quad (1.18)$$

Dabei ist wieder

$$x^{(k)}(0) = \left. \frac{d^k}{dt} x(t) \right|_{t=0}, \quad k = 1, 2, \dots$$

Nun ist $x_0^{(1)} = Ax_0$, $x_0^{(2)} = Ax_0^{(1)} = A^2 x_0$, $x_0^{(3)} = Ax_0^{(2)} = A^3 x_0$, etc. Setzt man diese Ausdrücke in (1.18) ein, so erhält man die Reihe

$$x(t) = x_0 \left(1 + tA + \frac{t^2}{2!} A^2 + \frac{t^3}{3!} A^3 + \dots \right) = x_0 \left(1 + At + \frac{(At)^2}{2!} + \frac{(At)^3}{3!} + \dots \right), \quad (1.19)$$

denn $tA = At$. In der Klammer rechts steht aber gerade die Exponentialmatrix von At , so dass man (1.17) erhält. $\square$

Jetzt kann die Lösung für das inhomogene System $\dot{x} = Ax + u$ bestimmt werden. Es gilt

Satz 1.2.2 *Gegeben sei das inhomogene Gleichungssystem $\dot{x} = Ax + u$ sowie einen Vektor x_0 von Anfangswerten. Dann sind die Lösungen $x(t)$ für dieses System durch*

$$x(t) = x_c(t) + x_p(t). \quad (1.20)$$

gegeben. Dabei ist

$$x_c(t) = e^{A(t-t_0)} x_0, \quad x_p(t) = \int_{t_0}^t e^{A(t-\tau)} u(\tau) d\tau. \quad (1.21)$$

mit $x_0 = x(t_0)$.

Beweis: Die Gleichung $\dot{x} = Ax + u$ impliziert $\dot{x} - Ax = u$. Multipliziert man diese Gleichung von links mit e^{-At} , so erhält man

$$e^{-At}\dot{x} - e^{-At}Ax = e^{-At}u$$

Sind f und g irgendzwei Funktionen von t , so ist die Ableitung des Produkts fg nach der Produktregel durch $(fg)' = f'g + fg'$ gegeben. Die vorangehende Gleichung legt nahe, dass sie über die Ableitung eines Produktes beschrieben werden kann. Es werde dazu die Ableitung des Produkts $e^{-At}x$ betrachtet; es ist

$$\frac{d}{dt}(e^{At}x) = e^{-At}\dot{x} - Ae^{-At}x = e^{-At}(\dot{x} - Ax) = e^{-At}u.$$

Integration bezüglich t liefert

$$\int \frac{d}{dt}(e^{At}x)dt = e^{At}x(t) = \int e^{At}u(t)dt.$$

Integriert man für die Grenzen t_0 (= Anfangszeitpunkt) und t , so erhält man insbesondere

$$\int_{t_0}^t \frac{d}{d\tau}(e^{A\tau}x)d\tau = e^{At}x(t) - e^{At_0}x(t_0) = \int_{t_0}^t e^{A\tau}u(\tau)d\tau$$

oder

$$e^{At}x(t) = e^{At_0}x(t_0) + \int_{t_0}^t e^{A\tau}u(\tau)d\tau$$

Multiplikation dieser Gleichung mit e^{-At} liefert dann

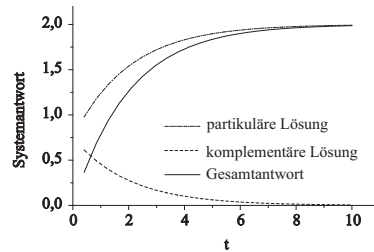
$$x(t) = e^{-A(t-t_0)}x(t_0) + \int_{t_0}^t e^{-A(t-\tau)}u(\tau)d\tau \quad (1.22)$$

und dies ist gerade (1.20). □

Bemerkungen:

1. Allgemein setzt sich die Lösung der linearen (affinen) Differentialgleichung aus zwei Teillösungen zusammen:
 - $x_c(t)$, die *komplementäre Lösung*
 - $x_p(t)$, die *partikuläre Lösung*
2. Der komplementäre Teil x_c der Lösung hängt von der "Auslenkung" des Systems zur Zeit t_0 ab, zu dem der Input (die "Störung") des Systems durch u beginnt, während x_p nur durch den Input u bestimmt wird.
3. Die Art der Reaktion des Systems hängt offenbar von der Exponentialmatrix $H(t) := e^{At}$ ab. Ist H bekannt, läßt sich die Reaktion des Systems auf einen beliebigen Input u berechnen. Um H weiter zu spezifizieren, wird im folgenden Teilabschnitt ein spezieller Input u betrachtet, der für die experimentelle Charakterisierung der Matrix $H(t) = e^{At}$ und damit des Systems von Bedeutung ist.

Abbildung 1.1: Gesamtantwort sowie die Komplementäre und partikuläre Antwort auf eine Schrittfunktion



Beispiel 1.2.1 Es sei $n = 1$ und das Eingangssignal sei eine Schrittfunktion, d.h. es sei

$$u(t) = \begin{cases} 0, & t \leq 0 \\ 1, & t > 0 \end{cases} \quad (1.23)$$

Für $n = 1$ ist $A = a$, $a \in \mathbb{R}$ eine einzelne reelle Zahl. Die Anwendung von (1.22) liefert, wenn man der Einfachheit $t_0 = 0$ annimmt,

$$x(t) = x(0)e^{-at} + \int_0^t e^{-a(t-\tau)} d\tau \quad (1.24)$$

Es ist aber

$$\int_0^t e^{-a(t-\tau)} d\tau = e^{-at} \int_0^t e^{a\tau} d\tau = e^{-at} \frac{1}{a} (e^{at} - 1) = \frac{1}{a} (1 - e^{-at}),$$

so dass sich schließlich

$$x(t) = x(0)e^{-at} - \frac{1}{a} (1 - e^{-at}) \quad (1.25)$$

ergibt. Für $x(0) = 0$ befindet sich das System in der Ruhelage, für $x(0) \neq 0$ ist das System bereits aus der Ruhelage ausgelenkt, wenn die Störung durch $u(t)$ beginnt.

Man muß jetzt zwei Fälle betrachten: (i) $a < 0$ und (ii) $a > 0$. Für den Fall $x(0) \neq 0$ strebt die komplementäre Lösung $x(0)e^{-at}$ gegen Null, wenn $a > 0$, und die partikuläre Lösung strebt gegen $1/a$. Für $a > 0$ strebt also die Antwort des Systems gegen die Konstante $1/a$; das System ist *stabil*. Für den Fall $a < 0$ wächst e^{-at} allerdings unbeschränkt; das System ist *instabil*. Der logisch mögliche Fall $a = 0$ ist aus offensichtlichen Gründen uninteressant.

□

1.2.4 Die Dirac-Impuls-Funktion und die Impulsantwort-Matrix

Definition 1.2.2 Die Funktion $\delta : \mathbb{R} \rightarrow \mathbb{R}$ sei durch die folgenden Bedingungen charakterisiert:

1.

$$\delta(\xi) = \begin{cases} \infty, & \xi = 0 \\ 0, & \xi \neq 0 \end{cases} \quad (1.26)$$

2.

$$\int_{-\infty}^{\infty} \delta(\xi) d\xi = 1 \quad (1.27)$$

3. Es sei f eine beliebige Funktion. Dann gilt

$$\int_{-\infty}^{\infty} f(\xi) \delta(x - \xi) d\xi = f(x) \quad (1.28)$$

Dann heit δ Dirac-Delta-Funktion oder Dirac-Impuls⁶.

Die δ -Funktion hat auf den ersten Blick bizarre Eigenschaften: sie ist nur fr $x = 0$ von Null verschieden; fr $x = 0$ nimmt sie den Wert ∞ an, und das Integral ber δ soll den endlichen Wert 1 haben. In der Tat ist die δ -Funktion keine gewhnliche Funktion, sondern eine *verallgemeinerte Funktion*. Man kann sich ihre Definition aber veranschaulichen.

Beispiel 1.2.2 Dazu betrachte man die Funktion

$$f(\xi) = \begin{cases} 0, & x < -a/2 \\ 1/a, & -a/2 \leq \xi \leq a/2 \\ 0, & \xi > a/2 \end{cases}$$

f definiert offenbar ein Rechteck der Breite a und der Hhe $1/a$, so dass fr alle $a \neq 0$

$$\int_{-\infty}^{\infty} f(\xi) d\xi = a \frac{1}{a} = 1$$

gilt. Fr $a \rightarrow 0$ strebt f dann gegen eine δ -Funktion. □

Beispiel 1.2.3 Die Gaufunktion

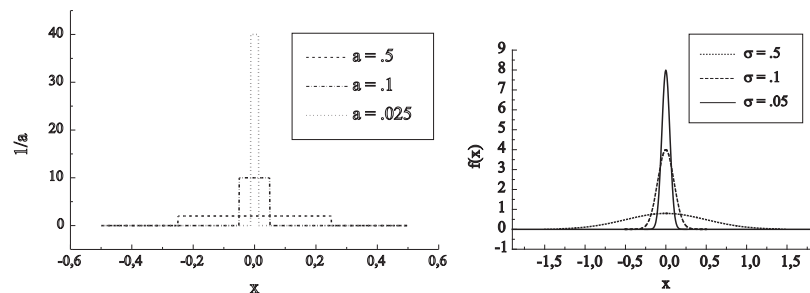
$$f(\xi) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{\xi^2}{2\sigma^2}\right)$$

ist ein weiteres Beispiel. Sicherlich ist das Integral ber f gleich 1 (f ist ja eine Dichtefunktion), und fr $\sigma \rightarrow 0$ strebt f gegen eine δ -Funktion. □

Die Bedingung (1.28) ist intuitiv einleuchtend: δ nimmt ja nur dann einen von Null verschiedenen Wert an, wenn das Argument ξ gleich Null ist. Dies ist der Fall, wenn $\xi = x$ ist. Gleichwohl ist die Forderung (1.28) nicht trivial, denn fr diesen Fall ist ja δ gleich unendlich! dass sie Sinn macht, kann man anhand der in den

⁶Nach ihrem Erfinder, dem Physiker Paul Adrien Maurice Dirac, (1902 – 1984), Mathematiker, wesentliche Leistungen in der Quantenmechanik.

Abbildung 1.2: Approximation der Dirac-Funktion: Rechteckpulse und Gauß-Funktionen



Beispielen 1.2.2 und 1.2.3 betrachteten Funktionen überprüfen, wenn man a bzw. σ na egration gegen Null gehen läßt.

Wählt man nun für u einen Vektor, dessen Komponenten aus δ -Funktionen $\delta(\tau)$ bestehen, so liefert insbesondere (1.28) die Beziehung

$$\int_{-\infty}^{\infty} e^{A(t-\tau)} u(\tau) d\tau = e^{At} \quad (1.29)$$

Dementsprechend heißt $H(t) = e^{At}$ auch die *Impulsantwortmatrix*.

Das System ist durch die Matrix A der Kopplungskoeffizienten bestimmt. Die Impulsantwortmatrix H ist wiederum durch A bestimmt. Gibt man nun einen Impuls (im Sinne der δ -Funktion) auf das System, so ist die Antwort gerade durch $H(t)$ gegeben. Gelingt es, $H(t)$ aus der Beobachtung abzuschätzen und kann A aus H berechnet werden, so hat man damit das System *identifiziert*. Auf die tatsächliche Durchführung dieses Ansatzes wird später zurückgekommen.

1.2.5 Die Operator-Schreibweise

Der lineare Operator und das Superpositionsprinzip Für die hier verfolgten Zwecke kann man sagen, daß ein System führt ein Inputsignal in eine Antwort bzw. in ein Outputsignal überführt. Dies läßt sich durch die Operatorschreibweise ausdrücken. Dabei wird das System durch den *Operator* L repräsentiert; es gilt

$$L(u) = x \quad (1.30)$$

wobei u und x natürlich (vektorielle) Funktionen etwa von t sind, also $u = u(t)$ und $x = x(t)$. Man zeigt nun leicht die Gültigkeit von

Satz 1.2.3 *Der Operator L (bzw. das durch den Operator L repräsentierte System) sei linear. Dann gilt*

$$L(\alpha_1 u_1 + \alpha_2 u_2) = \alpha_1 x_1 + \alpha_2 x_2 \quad (1.31)$$

wenn $L(u_1) = x_1$ und $L(u_2) = x_2$.

Beweis: Aus $\dot{x}_i = Ax_i + u_i$, $i = 1, 2$ folgt $u_i = \dot{x}_i - Ax_i$ und damit $\alpha_i u_i = \alpha_i \dot{x}_i - \alpha_i Ax_i$ und sicherlich

$$\alpha_1 u_1 + \alpha_2 u_2 = \alpha_1 \dot{x}_1 + \alpha_2 \dot{x}_2 - \alpha_1 Ax_1 - \alpha_2 Ax_2 = \alpha_1 \dot{x}_1 + \alpha_2 \dot{x}_2 - A(\alpha_1 x_1 + \alpha_2 x_2)$$

bzw.

$$\alpha_1 \dot{x}_1 + \alpha_2 \dot{x}_2 = A(\alpha_1 x_1 + \alpha_2 x_2) + \alpha_1 u_1 + \alpha_2 u_2$$

und dies ist gerade die Aussage des Satzes: man muß nur $x = \alpha_1 x_1 + \alpha_2 x_2$ setzen. $\square$

Bemerkungen:

1. Die Aussage (1.31) heißt auch *Superpositionsprinzip*.
2. Das Superpositionsprinzip kann auch zur Definition von linearen Systemen gewählt werden. Während das Superpositionsprinzip allerdings leicht aus der linearen Differentialgleichung $\dot{x} = Ax + u$ folgt, ist es schwieriger, die Differentialgleichung aus dem Superpositionsprinzip herzuleiten. Deshalb wurden hier lineare Systeme durch die linearen Differentialgleichungen definiert.

Die Impulsantwort und die Darstellung von x als Faltungsintegral Es sei $\Delta(t) = (\delta_1(t), \delta_2(t), \dots, \delta_n(t))'$ und das Inputsignal sei durch Δ definiert, d.h. es möge $u(t) = \Delta(t)$ gelten. Dann kann man

$$L(u) = L(\Delta) = H(t) \quad (1.32)$$

schreiben; die Antwort des Systems ist gerade die Impulsantwort. Man bemerke, dass die Impulsantwort des Systems durch (1.32) auch *definiert* werden kann; wählt man diesen Weg, hat man allerdings die Aufgabe, zu zeigen, dass H gerade die Form e^{At} hat.

Der Nutzen von (1.32) besteht andererseits u.a. darin, dass Eingangs- und Ausgangssignal stets durch eine *Faltung* aufeinander bezogen sind:

Definition 1.2.3 *Es seien f und h irgendzwei Funktionen. Dann heißt das Integral*

$$x(t) = \int_{-\infty}^{\infty} f(\tau)h(t - \tau)d\tau \quad (1.33)$$

ein Faltungsintegral oder einfach Faltung von f und h . Man schreibt auch abkürzend

$$x = f * h \quad (1.34)$$

Nach Satz 1.2.2, p. 15, ist die partikuläre Lösung der Gleichung $\dot{x} = Ax + u$, d.h. der Teil der Lösung, der durch den Input u "erzwungen" wird, durch

$$x_p(t) = \int_{t_0}^t e^{-A(t-\tau)}u(\tau)d\tau$$

gegeben. Das Integral auf der rechten Seite hat die Form einer *Faltung*, auch wenn $e^{-A(t-\tau)}u(\tau)$ ein Vektor ist. Generell gilt der folgende Satz:

Satz 1.2.4 *Gilt die Beziehung (1.33), so gilt auch*

$$x(t) = \int_{-\infty}^{\infty} f(t-\tau)h(\tau)d\tau \quad (1.35)$$

Beweis: Der Beweis erfolgt durch Substitution von $\xi = t - \tau$ und nachfolgende Umbenennung der Variablen. $\square$

Der Ausdruck (1.29) für x_p ist eine Vektorgleichung. Ihre Struktur wird deutlicher, wenn man sie etwas ausführlicher anschreibt. Es seien $h_{ij}(t)$ die Elemente der Matrix $H(t) = e^{At}$. Die i -te Komponente des Vektors $e^{A(t-\tau)}u(\tau)$ ist dann durch das Skalarprodukt der i -ten Zeile von $e^{A(t-\tau)}$ mit dem Vektor $u(\tau)$ gegeben, d.h. durch

$$h_{i1}(t-\tau)u_1(\tau) + \cdots + h_{in}(t-\tau)u_n(\tau) = \sum_{j=1}^n h_{ij}(t-\tau)u_j(\tau),$$

u_j die j -te Komponente von u , und die i -te Komponente des Antwortvektors x_{ip} ist durch

$$x_{ip}(t) = \int_{t_0}^t \sum_{j=1}^n h_{ij}(t-\tau)u_j(\tau)d\tau = \sum_{j=1}^n \int_{t_0}^t h_{ij}(t-\tau)u_j(\tau)d\tau \quad (1.36)$$

gegeben (Vertauschbarkeit von Summation und Integration vorausgesetzt). Man sieht, dass die Komponenten des Antwortvektors durch eine Summe von einfachen (nicht vektoriellen) Faltungsintegralen gegeben sind. Sind die Komponenten von u insbesondere δ -Funktionen, so erhält man

$$x_{ip}(t) = \sum_{j=1}^n \int_{t_0}^t h_{ij}(t-\tau)\delta(\tau)d\tau = \sum_{j=1}^n h_{ij}(t). \quad (1.37)$$

Beispiel 1.2.4 Ein besonders einfaches Beispiel ist der Fall $n = 1$. Das System $\dot{x} = Ax + u$ geht dann in das System $\dot{x} = ax + u$ über, wobei $a \in \mathbb{R}$, denn die Matrix A hat jetzt nur ein Element, nämlich a . Die Impulsantwort des Systems ist dann

$$H(t) = ce^{at}, \quad a \in \mathbb{R}, \quad c \in \mathbb{R} \quad (1.38)$$

wobei c eine Proportionalitätskonstante ist. Statt (1.38) kann dann auch

$$L[\delta(t)] = c \int_0^t e^{a(t-\tau)}\delta(\tau)d\tau = ce^{at}$$

geschrieben werden. Für $a > 0$ wächst die Antwort als Reaktion auf den Impuls offenbar unbeschränkt (zumindest so lange, wie die Rahmenbedingungen, unter denen das System definiert ist, gelten); das System ist *instabil*. Für $a < 0$ fällt die Antwort von ihrem Maximum c exponentiell gegen Null ab (das System *relaxiert*); das System ist *stabil*. $\square$

Die in Beispiel 1.38 angesprochene Frage der Stabilität eines linearen Systems kann nur diskutiert werden, wenn bekannt ist, wie $H(t) = e^{At}$ im allgemeinen Fall $n > 1$ berechnet wird. Darauf wird im folgenden Abschnitt eingegangen.

1.2.6 Zur Berechnung der Impulsmatrix

A ist eine $n \times n$ -Matrix. Es sei T eine $n \times n$ -Matrix, für die die inverse Matrix T^{-1} existieren möge, so dass $TT^{-1} = T^{-1}T = I$, I die Einheitsmatrix. Weiter sei $T^{-1}AT = \Lambda$, Λ wieder eine $n \times n$ -Matrix, so dass $A = T\Lambda T^{-1}$. Insbesondere sei Λ die Diagonalmatrix der Eigenwerte von A , wobei vorausgesetzt werde, dass alle Eigenwerte von A verschieden seien. Die Matrix T enthält dann die Eigenvektoren von A . Ist ein Eigenwert λ_i komplex, so dass $\lambda_i = \lambda_{i1} + j\lambda_{i2}$ mit $j = \sqrt{-1}$ gilt, so existiert stets ein zweiter Eigenwert $\lambda_k = \bar{\lambda}_i = \lambda_{i1} - j\lambda_{i2}$ und die zugehörigen Eigenvektoren sind ebenfalls konjugiert komplex (für eine Begründung dieser Aussage lese man in Lehrbüchern der linearen Algebra nach). Es gilt nun der

Satz 1.2.5 *Es sei T die Matrix der Eigenvektoren von A , und Λ sei die zugehörige Diagonalmatrix der Eigenwerte von A . Dann gilt*

$$e^{At} = Te^{\Lambda t}T^{-1}, \quad (1.39)$$

wobei

$$e^{\Lambda t} = \begin{pmatrix} e^{\lambda_1 t} & 0 & \dots & 0 \\ 0 & e^{\lambda_2 t} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & e^{\lambda_n t} \end{pmatrix} \quad (1.40)$$

ist.

Beweis: Aus $T^{-1}AT = \Lambda$ folgt sofort $A = T\Lambda T^{-1}$. Weiter ist $A^2 = (T\Lambda T^{-1})^2 = T\Lambda T^{-1}T\Lambda T^{-1} = T\Lambda^2 T^{-1}$, $A^3 = A^2A = T\Lambda^2 T^{-1}T\Lambda T^{-1} = T\Lambda^3 T^{-1}$, etc, so dass man allgemein

$$A^k = T\Lambda^k T^{-1} \quad (1.41)$$

folgern kann. Analog zeigt man, dass

$$(At)^k = T(\Lambda t)^k T^{-1} \quad (1.42)$$

Es ist aber

$$e^{At} = \sum_{k=0}^{\infty} \frac{(At)^k}{k!} = \sum_{k=0}^{\infty} \frac{T(\Lambda t)^k T^{-1}}{k!} = T \left(\sum_{k=0}^{\infty} \frac{(\Lambda t)^k}{k!} \right) T^{-1} = Te^{\Lambda t}T^{-1}$$

womit (1.39) gezeigt ist. □

(1.39) zeigt, dass man zur Berechnung von e^{At} nicht alle Terme $(At)^k/k!$ der unendlichen Reihe (1.12) berechnen muß. Es genügt vielmehr, die Eigenwerte und die

zugehörigen Eigenvektoren der Matrix A zu berechnen, um e^{At} für alle t bestimmen zu können.

Insbesondere kann man nun einen expliziten Ausdruck für die Elemente $h_{ij}(t)$ der Impulsantwortmatrix angeben. Es sei t_{ij} das Element in der i -ten Reihe und j -ten Spalte von T (d.h. t_{ij} sei die i -te Komponente des j -ten Eigenvektors von A), und $\tilde{t}_{ij}$ sei das entsprechende Element von T^{-1} . Dann gilt wegen $H = e^{At} = T e^{\Lambda t} T^{-1}$

$$h_{ij}(t) = \sum_{k=1}^n t_{ik} \tilde{t}_{kj} e^{\lambda_k t} = \sum_{k=1}^n t_{ik} \tilde{t}_{kj} e^{\lambda_k t} \quad (1.43)$$

Das Element h_{ij} ist also eine Summe von Exponentialfunktionen $e^{\lambda_k t}$, "gewichtet" mit den $t_{ik} \tilde{t}_{kj}$. Auf die inhaltliche Bedeutung dieses Sachverhalts wird eingegangen, wenn die verschiedenen Formen der Zusammenschaltung von Systemen diskutiert wird.

1.2.7 Stabilität linearer Systeme; Eigenfrequenzen

Nach (1.37) ist nun $x_{pi}(t) = \sum_j h_{ij}(t)$, und damit ist

$$x_{pi}(t) = \sum_{j=1}^n \sum_{k=1}^n t_{ik} \tilde{t}_{kj} e^{\lambda_k t} = \sum_{k=1}^n \left(\sum_{j=1}^n t_{ik} \tilde{t}_{kj} \right) e^{\lambda_k t} = \sum_{k=1}^n c_{ik} e^{\lambda_k t} \quad (1.44)$$

mit $c_{ik} = \sum_j t_{ik} \tilde{t}_{kj}$. Die i -te Komponente der Reaktion auf einen Impuls ist also eine Summe von Exponentialfunktionen. Hieraus wird sofort die Bedeutung der Eigenwerte von A deutlich. So sei $\lambda_k \in \mathbb{R}$ und weiter gelte $\lambda_k > 0$. Dann wächst $e^{\lambda_k t}$ unbeschränkt und damit auch $x_{pi}(t)$, d.h. das System ist *instabil*. Es sei weiter $\lambda_k \in \mathbb{C}$, wobei mit $\mathbb{C}$ die Menge der komplexen Zahlen bezeichnet wird, so dass $\lambda_k = \alpha_k + i\beta_k$. Dann existiert stets ein Eigenwert $\lambda_s = \bar{\lambda}_k = \alpha_k - i\beta_k$ und die Koeffizienten c_{ik} und c_{is} sind ebenfalls zueinander konjugiert komplex, d.h. $c_{is} = \bar{c}_{ik} = |c_{ik}| e^{-i\phi}$, $\bar{c}_{ik} = |c_{ik}| e^{i\phi}$, so dass wir für diese beiden Eigenwerte die Teilsumme

$$S_{ik}(t) := |c_{ik}| e^{i\phi} e^{\lambda_k t} + |c_{ik}| e^{-i\phi} e^{\bar{\lambda}_k t} = |c_{ik}| e^{\alpha_k t} \left(e^{i(\beta_k t + \phi)} + e^{-i(\beta_k t + \phi)} \right)$$

d.h.

$$S_{ik}(t) = |c_{ik}| e^{i\phi} e^{\lambda_k t} + |c_{ik}| e^{-i\phi} e^{\bar{\lambda}_k t} = 2e^{\alpha_k t} |c_{ik}| \cos(\beta_k t + \phi). \quad (1.45)$$

Die Existenz eines Paares konjugiert komplexer Eigenwerte bedeutet also, dass die entsprechende Komponente der Impulsantwort mit der Frequenz β_k schwingt; dies ist eine *Eigenschwingung* des Systems, und β_k heißt dementsprechend eine *Eigenfrequenz* des Systems. Gibt es mehrere Paare von Eigenwerten, so gibt es verschiedene Eigenfrequenzen und damit verschiedene Eigenschwingungen, die sich additiv überlagern.

Aus (1.45) folgt sofort eine Aussage über die Stabilität von linearen Systemen. Für $|c_{ik}| < \infty$ ist auch $|c_{ik}| \cos(\beta_k t + \phi) < \infty$ für alle t . $S_{ik}(t)$ hängt allerdings noch von dem Faktor $e^{\alpha_k t}$ ab. α_k ist der Realteil des Eigenwerts λ_k . Für $\alpha_k > 0$ strebt

dieser Faktor gegen unendlich, damit auch $\Sigma_{ik}(t)$ und damit die Systemantwort insgesamt: das System ist instabil. Für $\alpha_k < 0$ gegen Null, und damit $S_{ik}(t)$. Damit gilt

Satz 1.2.6 *Das System $\dot{x} = Ax + u$ ist stabil, wenn die Realteile aller Eigenwerte λ_k von A negativ sind.*

Definition 1.2.4 *Die Funktionen $\exp(\lambda_j t)$ heißen die Moden des Systems, und die Vektoren $e^{\lambda_j t} T_j$ heißen die dynamischen Moden des Systems. Die Matrix der T der Eigenvektoren heißt auch Modalmatrix.*

1.2.8 Anmerkungen zum Zustandsbegriff

Für das System $\dot{x} = Ax + u$ - die folgenden Betrachtungen übertragen sich auf das verallgemeinerte System (1.7) - waren die Komponenten des Vektors x als Zustandsvariablen bezeichnet worden. Die Definition des Zustands eines linearen Systems ist aber nicht eindeutig. So sei etwa T eine $n \times n$ Matrix mit konstanten Koeffizienten, für die die Inverse T^{-1} existieren möge (T sei also nichtsingulär). Ist x nun ein Zustandsvektor, so kann man durch die Transformation $T^{-1}x$ zu einem neuen n -dimensionalen Vektor $z = z(t)$ übergehen. Für jeden Zeitpunkt t korrespondiert dann zu x der Vektor z . Multipliziert man von links mit T , so erhält man $TT^{-1}x = x = Tz$, so dass man auch sagen kann, der Vektor x sei durch Transformation aus einem Vektor z hervorgegangen. Da T eine Matrix mit konstanten Elementen ist, folgt $\dot{x} = T\dot{z}$, und die Gleichung

$$\dot{x} = Ax + u$$

kann dann in der Form

$$T\dot{z} = ATz + u$$

geschrieben werden. Multipliziert man nun von links mit T^{-1} , so erhält man die Gleichung

$$\dot{z} = T^{-1}ATz + T^{-1}u. \quad (1.46)$$

Für den Zustandsvektor z wird das System, d.h. werden die Kopplungen zwischen den Zustandsvariablen, also durch die Matrix $T^{-1}AT$ und nicht mehr durch die Matrix A charakterisiert. Natürlich ergibt sich jetzt die Frage, welche Aspekte der Systembeschreibung denn invariant sind, wenn schon die Definition der Zustandsvariablen nicht eindeutig ist. Es zeigt sich, dass die Eigenwerte von A und $T^{-1}AT$ die gleichen sind. Die Eigenwerte von A ergeben sich als Wurzeln des Polynoms $Q(\lambda) = |A - \lambda I| = 0$. Für die Matrix $T^{-1}AT$ hat man für die Eigenwerte ein Polynom $\tilde{Q}(\mu) = |T^{-1}AT - \mu I| = 0$. Nun gilt aber $I = T^{-1}T$, so dass man auch $\tilde{Q}(\mu) = |T^{-1}AT - \mu T^{-1}T| = 0$ schreiben kann. Dies ist identisch mit $\tilde{Q}(\mu) = |T^{-1}(AT - \mu T)| = |T^{-1}(A - \mu I)T|$. Aus der Matrixalgebra ist aber bekannt, dass die Determinante eines Produktes von Matrizen gleich dem Produkt der entsprechenden Determinanten ist, also folgt $\tilde{Q}(\mu) = |T^{-1}||A - \mu I||T| = 0$. Die Determinanten $|T^{-1}|$ und $|T|$ sind nach Voraussetzung ungleich Null, sonst wäre die

Matrix T nicht nichtsingulär. Also muß $|A - \mu I| = 0$ gelten, d.h. aber $|A - \mu I| = Q(\mu)$, so dass die Eigenwerte μ von $T^{-1}AT$ gerade gleich den Eigenwerten λ von A sein müssen. Die Eigenwerte von A oder $T^{-1}AT$ beschreiben aber gerade die qualitativ wichtigen Aspekte der Dynamik des Systems, wie etwa Stabilität in Gleichgewichtslagen und die Eigenfrequenzen für den Fall komplexer Eigenwerte.

Es werde nun insbesondere angenommen, dass die Eigenwerte alle verschieden sind und dass T die Modalmatrix ist; dann gilt $T^{-1}AT = \Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$, und nach (1.46) erhält man

$$\dot{z} = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{pmatrix} z + T^{-1}u \quad (1.47)$$

Betrachtet man als Ausgangsvariable beim ursprünglichen System $y = Cx + Du$, so ergibt sich jetzt für y die Gleichung $y = \tilde{C}z + Du$, wobei $\tilde{C} = CT$ ist. Das Wesentliche an (1.47) ist aber, dass die Zustandsvariablen z_i nun *entkoppelt* sind. Wie später noch gezeigt werden wird, bedeutet dies, dass sich das System als *Parallelschaltung* bestimmter elementarer Teilsysteme auffassen läßt, weil ja die Elemente der Matrix Λ die Kopplungskoeffizienten der Zustandsvariablen enthält, und hier $\lambda_{ij} = 0$ für $i \neq j$ ist.

Der Fall mehrerer gleicher Eigenwerte ist komplizierter und soll hier nicht eingehend diskutiert werden; allgemein läßt sich sagen, dass er sich ergibt, wenn das System aus parallel arbeitenden Teilsystemen aufbaubar ist, die ihrerseits als rückwirkungsfreie Hintereinander- oder Reihenschaltung gleichartiger Subsysteme repräsentiert werden können.

1.2.9 Lineare Systeme n -ter Ordnung

Systeme der Form (1.3) bzw. (1.4), p. 1.4, sind eine allgemeine Repräsentation von linearen Systemen. Die Diskussion bestimmter dynamischer Systeme kann allerdings auf Differentialgleichungen führen, die Ableitungen höherer Ordnung enthalten. Es wird die folgende Notation eingeführt: mit D werde der *Differentiationsoperator* bezeichnet. Ist $y(t)$ eine mindestens n -mal differenzierbare Funktion von t , so bezeichne Dy das Differential $dy(t)/dt$, D^2y die entsprechende zweite Ableitung $d^2y(t)/dt^2$, allgemein $D^n y$ dann die n -te Ableitung $d^n y(t)/dt^n$.

Beispiel 1.2.5 (Harmonischer Oszillator) Es sei ein System gegeben, das um den Betrag $z(t)$ aus seiner Ruhelage ausgelenkt wird. Je mehr sie ausgelenkt wird, desto größer sei die "Rückstellkraft" $F = F(z)$, mit der es in den Ausgangszustand zurückzugelangen versucht (man denke an eine Feder). Man kann annehmen, dass diese Kraft der Auslenkung proportional ist⁷, und da sie der Auslenkung entgegengesetzt wirkt, kann man $z(t) = -kF(z)$ schreiben, wobei k eine Proportionalitätskonstante ist. Nun ist aber die Kraft proportional der Beschleunigung, und die ist

⁷Hookesches Gesetz

durch die zweite Ableitung der Auslenkung gegeben, also durch

$$D^2 z(t) = \ddot{z}(t) = \frac{dz^2(t)}{dt^2}.$$

Mithin ist die "Bewegungsgleichung" des Systems

$$z(t) = -k\ddot{z}(t), \quad \text{bzw. } kD^2 z(t) + z(t) = 0. \quad (1.48)$$

Die DGL ist linear. Da $z(t)$ proportional seiner zweiten Ableitung ist, muß es sich bei $z(t)$ um eine Funktion wie $\exp(\lambda t)$ oder $\sin(\omega t + \phi)$ bzw. um lineare Kombinationen solcher Funktionen handeln. Da die Sinusfunktion selbst bereits die Summe zweier Exponentialfunktionen ist $\sin(\omega t) = (e^{j\omega t} - e^{-j\omega t})/2j$, kann man gleich die Hypothese $z(t) = \alpha \sin(\omega t + \phi)$ betrachten, wobei α eine Konstante ist. Dann ist $\ddot{z}(t) = -\alpha\omega^2 \sin(\omega t + \phi)$. In (1.48) eingesetzt ergibt sich sofort $(-k\omega^2 + 1) \sin(\omega t + \phi) = 0$, woraus dann $\omega = 1/\sqrt{k}$ folgt. Einmal angestoßen schwingt das System also sinusförmig mit der Frequenz $\omega = 1/\sqrt{k}$. Die Frequenz ω ist eine Eigenfrequenz und durch das Material der Feder bestimmt.

Das ungedämpfte Schwingen ist eine Konsequenz der Annahme $z(t) = -k\ddot{z}(t)$. Um der tatsächlich auftretenden Dämpfung gerecht zu werden (eine Feder hört irgendwann auf, zu schwingen), muß die innere Reibung des Materials bei der Bewegung berücksichtigt werden. Eine vernünftige Annahme ist, dass diese Reibung proportional zur Geschwindigkeit der Bewegung, d.h. zu $\dot{z}(t) = Dz(t) = dz(t)/dt$ ist, und dass sie additiv auf die Bewegung wirkt. Dann hat man die Dgl

$$a_2 D^2 z(t) + a_1 Dz(t) + a_0 z(t) = 0. \quad (1.49)$$

Man kann vermuten, dass die Lösung $z(t)$ wieder sinusförmig schwingt. Der direkte Ansatz

$$z(t) = \alpha \sin(\omega t + \phi)$$

ist aber nicht so einfach zu handhaben, da die erste Ableitung $\dot{z}(t) = \alpha\omega \cos(\omega t + \phi)$ auftaucht. Man beginnt am besten mit dem Ansatz $z(t) = e^{\lambda t}$ und findet

$$(a_2 \lambda^2 + a_1 \lambda + a_0) e^{\lambda t} = 0,$$

d.h.

$$a_2 \lambda^2 + a_1 \lambda + a_0 = 0. \quad (1.50)$$

da ja $e^{\lambda t} \neq 0$ für $t < \infty$. Dies ist das *charakteristische Polynom* der Dgl, das hier insbesondere eine quadratische Gleichung in λ ist, die bekanntlich die Lösungen

$$\lambda_{1,2} = -\frac{a_1}{2a_2} \pm \frac{1}{2a_2} \sqrt{a_1^2 - 4a_0a_2} \quad (1.51)$$

hat. Man hat die folgenden Fälle:

- a. **Aperiodischer Grenzfall.** Für $a_1^2 = 4a_0a_2$ verschwindet die Wurzel in (1.51) und es gibt nur eine reelle Lösung $\lambda = -a_1/(2a_2)$. Es werde nun $a_1 \neq 0$ und $a_2 \neq 0$ vorausgesetzt. Offenbar ist $\lambda < 0$, wenn a_1 und a_2 beide das gleiche

Vorzeichen haben; dann ist die Lösung stabil. Für ungleiches Vorzeichen ist $\lambda > 0$ und die Lösung ist instabil. Aber mit $e^{\lambda t}$ ist auch $y(t) = te^{\lambda t}$ eine Lösung: $\dot{y}(t) = e^{\lambda t} + \lambda te^{\lambda t} = (1 + \lambda t)e^{\lambda t}$, $\ddot{y}(t) = (2\lambda + t\lambda^2)e^{\lambda t}$, so dass, unter Berücksichtigung von $\lambda = -a_1/(2a_2)$, die DGL erfüllt ist. Sind α_1 und α_2 beliebige reelle Zahlen, so ist damit die Funktion $z(t) = (\alpha_1 + \alpha_2 t)e^{\lambda t}$ die allgemeine Lösung. Diese Lösung repräsentiert den *aperiodischen Grenzfall*.

Es sei $a_1 = 0$. Da $a_1^2 = 4a_0a_2$ vorausgesetzt wurde, folgt, dass dann auch $a_0 = 0$ oder $a_2 = 0$ oder $a_0 = a_2 = 0$ gelten muß. Der Vergleich mit (1.48) und (1.49) zeigt, dass dann die Voraussetzung, dass es sich überhaupt um ein sich bewegendes System handeln soll, nicht mehr erfüllt sind.

- b. **Kriechfall.** Nun sei $a_1^2 > 4a_0a_2$. Dann ist die Wurzel reell und damit sind die beiden Werte λ_1 und λ_2 reell. Die allgemeine Lösung ist $z(t) = \alpha_1 e^{\lambda_1 t} + \alpha_2 e^{\lambda_2 t}$. Betrachten wir noch den Fall $z(t) = 0$: dann ist $\alpha_1/\alpha_2 = e^{(\lambda_1 - \lambda_2)t}$; da die Exponentialfunktion monoton in t ist, kann es höchstens einen Wert $t = t_0$ geben, für den diese Bedingung erfüllt ist, d.h. die Lösung hat höchstens einen Nulldurchgang. Im stabilen Fall muß weiter $z(t) \rightarrow 0$ für $t \rightarrow \infty$ gelten. Man spricht vom *Kriechfall*.
- c. **Oszillatorischer Fall** Der letzte Fall ist $a_1^2 < 4a_0a_2$. Dann ist die Wurzel durch $\pm j\sqrt{4a_0a_2 - a_1^2}$, $j = \sqrt{-1}$, gegeben; setzt man $\mu = \sqrt{4a_0a_2 - a_1^2}$, ist $\lambda_1 = -a_1/2a_2 + j\mu$, $\lambda_2 = -a_1/2a_2 - j\mu$, und die allgemeine Lösung ist

$$z(t) = e^{(-a_1/2a_2)t}(\alpha_1 e^{j\mu t} + \alpha_2 e^{-j\mu t}).$$

Sind α_1 und α_2 beliebige reelle Zahlen, so ist die Lösung offenbar komplex. Um eine physikalisch realisierbare Lösung zu erhalten, muß man Randbedingungen bezüglich α_1 und α_2 einführen. Die einfachste ist, $\alpha_1 \in \mathbb{C}$ zu wählen und $\alpha_2 = \bar{\alpha}_1$ als dazu konjugiert komplexe Zahl, so dass $\alpha_1 = |\alpha_1|e^{j\varphi}$, $\alpha_2 = |\alpha_1|e^{-j\varphi}$. Dann erhält man insbesondere

$$z(t) = e^{(-a_1/2a_2)t}|\alpha_1|\cos(\mu t + \varphi); \quad (1.52)$$

die Lösung oszilliert.

Es sei insbesondere $a_1 = 0$. Dann ist $e^{(-a_1/2a_2)t} \equiv 1$ und das System schwingt ungedämpft weiter für alle t , zumindest so lange, wie die Randbedingungen (die Werte von a_0 , a_1 und a_2) nicht geändert werden.

Die Lösungen sind hier gewissermaßen durch Raten gefunden worden, wobei das Raten natürlich durch das Wissen geleitet wurde, dass Exponentialfunktionen sich bei Differentiation reproduzieren. Andererseits ist in Satz 1.1.2 eine allgemeine Lösung für lineare Systeme angegeben worden, und daher ergibt sich die Frage, wie diese allgemeine Lösung auf die hier betrachtete Dgl zweiter Ordnung angewendet werden kann.

Dazu muß ein Zustandsvektor $x(t)$ definiert werden, der eben den Zustand des Oszillators in jedem Zeitpunkt t charakterisiert. Eine Möglichkeit besteht darin, den Zustand der Feder im Zeitpunkt t als (i) durch die Auslenkung $z(t)$, und (ii) durch

die Werte der Ableitungen $\dot{z}(t)$ und $\ddot{z}(t)$ gegeben anzunehmen. Dementsprechend kann man nun $x_1(t) = z(t)$, $x_2(t) = \dot{z}(t)$ setzen und hat damit $x(t) = (x_1(t), x_2(t))'$ definiert. Offenbar ist $\dot{x}(t) = (\dot{x}_1(t), \dot{x}_2(t))' = (\dot{z}(t), \ddot{z}(t))'$. Bedenkt man, dass die Dgl in der Form

$$\ddot{z}(t) = -\frac{a_1}{a_2}\dot{z}(t) - \frac{a_0}{a_2}z(t)$$

geschrieben werden kann, so hat man das homogene System

$$\dot{x}(t) = Ax(t) \quad (1.53)$$

mit der Koeffizientenmatrix $A = (a_{ij})$. Es werde nun die komplementäre Lösung $x_c(t) = x_0 \exp(At)$ betrachtet, wobei x_0 ein Startvektor ungleich Null sei. Dann muß $\exp(At)$ bestimmt werden. Dazu benötigt man die Eigenwerte von A , d.h. die Wurzeln des Polynoms

$$Q(\lambda) = |A - \lambda I| = \lambda^2 + \frac{a_1}{a_2}\lambda + \frac{a_0}{a_2} = 0,$$

die bekanntlich durch

$$\lambda_{1,2} = -\frac{1}{2}\frac{a_1}{a_2} \pm \frac{1}{2a_2}\sqrt{a_1^2 - 4a_0a_2}$$

gegeben sind. Die Elemente h_{ij} von $H = e^{At}$ sind durch Summen der Form $\alpha_1 e^{\lambda_1 t} + \alpha_2 e^{\lambda_2 t}$ definiert, und damit wird man zur gleichen Betrachtung wie oben geführt. Die Wahl der Zustandsvariablen $z(t)$, $\dot{z}(t)$ und $\ddot{z}(t)$ führt also auf die gleiche Lösung, die jetzt aber nach Satz 1.2.2, p. 15 gefunden wird. $\square$

Natürlich kann man noch den Fall betrachten, wo eine "Störung" $u(t)$ an den harmonischen Oszillator gelegt wird. Die Rechnung soll hier nicht im einzelnen durchgeführt werden, statt dessen soll gleich der allgemeine Fall betrachtet werden, bei dem nicht nur eine Ableitung zweiter Ordnung, sondern n -ter Ordnung mit $1 \leq n < \infty$ beliebig auftaucht. Man mag sich zuerst fragen, welche Bedeutung solche Ableitungen haben. Sind die Eigenwerte von A alle verschieden, so war bereits bei der Betrachtung von (1.47), p. 25, angemerkt worden, dass sich das System als Parallelschaltung von Teilsystemen auffassen läßt, die jeweils die Impulsantwort $\alpha_i e^{\lambda_i t}$ bzw. $\alpha_i \sin(\omega_i t + \phi_i)$ im Falle komplexer Eigenwerte haben. Es zeigt sich weiter, dass höhere Ableitungen auftreten, wenn im Gesamtsystem solche Teilsysteme rückwirkungsfrei hintereinandergeschaltet werden, wie in Abschnitt 1.2.12 näher ausgeführt wird. In jedem Fall kann man aber annehmen, dass es Zustandsvariablen $x(t) = (x_1(t), \dots, x_n(t))'$ gibt, mit denen das System in der Form $\dot{x} = Ax + u$, d.h. in der Form (1.3), p. 11, repräsentiert werden kann, da dieses ja die allgemeine Form der Repräsentation eines linearen Systems ist. Nun möchte man gelegentlich aber

nur die Beziehung zwischen einem bestimmten Zustand $x_i(t) = z(t)$ und einem Eingangssignal betrachten. Dann können alle anderen Variablen eliminiert werden: dadurch entsteht eine Dgl höherer Ordnung, bei der etwa n Ableitungen des Ausgangssignals z und $m \leq n$ Ableitungen der Eingangssignals u auftreten; man spricht dann von einer Dgl n -ter Ordnung. Sie hat die allgemeine Form

$$\sum_{k=0}^n a_k D^k z(t) = \sum_{j=1}^m b_j D^j u(t) \quad (1.54)$$

Die Entstehung von (1.54) werde an einem Beispiel verdeutlicht.

Beispiel 1.2.6 Gegeben sei das System $\dot{x} = Ax + u$, wobei $u = u(t)$ durch den Vektor $(b_1 w(t), b_2 w(t))'$ gegeben sei mit $b_1, b_2 \in \mathbb{R}$ und $w(t)$ eine (skalare) Funktion von t , die etwa den Verlauf eines Stimulus repräsentiere. Ausgeschrieben ergibt sich

$$\begin{aligned}\dot{x}_1 &= a_{11}x_1 + a_{12}x_2 + b_1 w \\ \dot{x}_2 &= a_{21}x_1 + a_{22}x_2 + b_2 w\end{aligned}$$

Es sei $z = x_2$ die gesuchte Funktion in Abhängigkeit von $w(t)$. Die Variable x_1 muß eliminiert werden. Die zweite Gleichung kann man nach x_1 auflösen:

$$x_1 = \frac{1}{a_{21}}(\dot{x}_2 - a_{22}x_2 - b_2 w)$$

Weiter kann man die zweite Gleichung noch einmal differenzieren; man erhält $\ddot{x}_2 = a_{21}\dot{x}_1 + a_{22}\dot{x}_2 + b_2\dot{w}$, was nach $\dot{x}_1$ aufgelöst

$$x_1 = \frac{1}{a_{21}}(\ddot{x}_2 - a_{22}\dot{x}_2 - b_2\dot{w})$$

ergibt. Diese Ausdrücke für x_1 und $\dot{x}_1$ können in die erste Gleichung eingesetzt werden. Nach einem Rearrangement der Terme hat man schließlich

$$\ddot{x}_2 - (a_{22} + a_{11})\dot{x}_2 + a_{21}(a_{11}a_{22} - a_{12})x_2 = a_{21}b_2\dot{w} + a_{21}b_1w;$$

dies ist eine *lineare Dgl mit konstanten Koeffizienten 2-ter Ordnung*. □

Es sei eine Dgl n -ter Ordnung wie (1.54) gegeben. Die Frage ist, wie man jetzt einen Zustandsraum definieren kann. Es wird zunächst der Fall

$$D^n z + a_{n-1}D^{n-1}z + \cdots + a_1 Dz + a_0 z = u \quad (1.55)$$

betrachtet; es treten also keine Ableitungen des Eingangssignales $u = u(t)$ auf. Die n -te Ableitung von z hat den Faktor 1, weil man die Gleichung stets durch den Faktor a_n dividieren kann; die entstehenden Quotienten sind hier wieder in a_{n-1}, a_{n-2} , etc. umbenannt worden. Wir definieren nun bestimmte Zustandsvariablen, die im gegebenen Zusammenhang auch *Phasenvariablen* genannt werden: es sei $x(t) = (x_1(t), \dots, x_n(t))'$ mit $x_1(t) = z(t)$, $x_2(t) = Dx_1(t)$, $x_3(t) = Dx_2(t) = D^2z(t)$, $\dots$, $x_n(t) = Dx_{n-1}(t) = D^n z(t)$. Setzt man diese Ausdrücke in (1.55) ein, so erhält man die Gleichung

$$\dot{x}_n + a_{n-1}x_n + a_{n-2}x_{n-1} + \cdots + a_1 x_2 + a_0 x_1 = u. \quad (1.56)$$

Die linke Seite kann als Skalarprodukt des Vektors

$$(1, a_{n-1}, \dots, a_0)'$$

mit dem Vektor

$$(\dot{x}_n, x_n, \dots, x_1)'$$

aufgefaßt werden. Um dieses Skalarprodukt in Form einer Matrixgleichung der Form (1.3) darzustellen, definiert man nun die Matrix A_c :

$$A_c = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 & 0 \\ 0 & 0 & 1 & \cdots & 0 & 0 \\ & & & \ddots & 1 & 0 \\ 0 & 0 & 0 & \cdots & 0 & 1 \\ -a_0 & -a_1 & -a_2 & \cdots & -a_{n-2} & -a_{n-1} \end{pmatrix}$$

Weiter sei $b_c = (0, 0, \dots, 0, 1)'$, $y(t) = x_1(t) = (1, 0, \dots, 0)'x = c'x$. Dann wird

$$\begin{aligned} x &= A_c x + b_c u \\ y &= c'x \end{aligned} \tag{1.57}$$

Damit ist die Dgl n -ter Ordnung als System von Dgln erster Ordnung dargestellt worden, auf das der Satz 1.1.2 angewandt werden kann, um die Lösung der Dgl zu finden.

Jetzt kann der allgemeine Fall betrachtet werden, wie er in (1.54) mit $m \neq 0$ gegeben ist, bei dem also Ableitungen des Eingangssignales $u(t)$ auftreten. (1.54) kann in der Form

$$(D^n + a_{n-1}D^{n-1} + \cdots + a_1D + a_0)z(t) = (b_mD^m + b_{n-1}D^{m-1} + \cdots + b_1D + b_0)u$$

geschrieben werden. Die Funktion $z(t)$ ist durch diese Dgl festgelegt. Man kann sich nun vorstellen, dass diese Funktion als *Eingangssignal* für ein System benutzt wird, dass nur durch die rechte Seite der Dgl definiert ist, d.h.

$$(b_mD^m + b_{n-1}D^{m-1} + \cdots + b_1D + b_0)x_1(t) = z(t);$$

hierdurch wird eine neue Funktion $x_1(t)$ definiert, eben als Lösung dieser Dgl. Gleichzeitig bedeutet aber diese Dgl gerade, dass für alle t aus dem betrachteten Intervall die Funktion $z(t)$ durch den Ausdruck $(b_mD^m + b_{n-1}D^{m-1} + \cdots + b_1D + b_0)x(t)$ gegeben ist, der also für $z(t)$ in die ursprüngliche Gleichung substituiert werden kann. Nun kann man ja $(b_mD^m + b_{n-1}D^{m-1} + \cdots + b_1D + b_0)$ als einen Operator auffassen, er werde mit $D_m(b)$ bezeichnet. Dann kann $D^n + a_{n-1}D^{n-1} + \cdots + a_1D + a_0$ durch $D_n(a)$ bezeichnet werden. Es gilt also

$$D_n(a)z(t) = u(t)$$

und mithin

$$D_n(a)D_m(b)x(t) = D_m(b)u(t).$$

Diese letzte Gleichung besagt, dass man sowohl auf $x(t)$ wie auf $u(t)$ den Operator $D_m(b)$ anwenden soll, und anschließend auf $D_m(b)x(t)$ noch einmal den Operator $D_n(a)$. Dann muß aber auch

$$D_n(a)x(t) = u(t)$$

gelten. Dies ist aber der bereits behandelte Fall, bei dem keine Ableitungen der Funktion $u(t)$ auftreten. Es gibt dann wieder eine Matrix A_c und einen Vektor b_c derart, dass

$$\dot{x}(t) = A_c x(t) + b_c u(t)$$

gilt, und das Ausgangssignal ist durch $z(t) = c'x$ gegeben, mit $c = (b_0, b_1, \dots, b_m, 0, \dots, 0)'$, für $m < n$; für $m = n$ treten in c keine Nullen auf. Damit ist gezeigt, dass sich allgemeine Dgln n -ter Ordnung auf Systeme von Dgln erster Ordnung zurückführen lassen.

$u(t)$ repräsentiere das Eingangssignal, und $g(t)$ sei die Zustandsvariable, die als Ausgangssignal gelten soll. Das betrachtete System sei durch eine Dgl der Form

$$\sum_{k=0}^n a_k g^{(k)}(t) = \sum_{j=1}^m b_j u^{(j)}(t) \quad (1.58)$$

beschreibbar. Ist speziell $u(t) = 0$ für alle t , so erhält man wie bei Systemen von Dgln erster Ordnung die *homogene Dgl*

$$\sum_{k=0}^n a_k g^{(k)}(t) = 0.$$

In Operatorschreibweise hat man dann

$$L[u(t)] = g(t).$$

Die Funktion $u(t)$ wird vorgegeben. Sind die Koeffizienten a_k und b_j bekannt, so kann man die Funktion durch Lösen der Dgl finden. Für den Fall, dass die Dgl nur erster Ordnung ist, ist die Lösung sofort auszurechnen (s. Beispiel 1.2.1, p. 17, Gleichung (1.24)); ist $u(t) = 0$ für $t < 0$, so erhält man für die partikuläre Lösung

$$\xi_p(t) = g(t) = \int_0^t e^{a(t-\tau)} u(\tau) d\tau = \int_0^t h(t-\tau) u(\tau) d\tau. \quad (1.59)$$

Dabei ist $h(t) = e^{at}$ die Impulsantwort des Systems; wählt man in der Tat $u(\tau) = \delta(\tau)$ so folgt aus (1.59) sofort $\xi_p(t) = g(t) = e^{at}$, d.h. aber, dass die Antwort auf $m\delta(\tau)$ durch me^{at} gegeben ist. Für Dgln erster Ordnung ergibt sich die partikuläre Lösung, d.h. die Antwort $g(t)$, also durch Faltung der Impulsantwort mit dem Eingangssignal. Bei *Systemen* von Dgln erster Ordnung ergibt sich die Lösung $\xi_p(t)$ ebenfalls durch Faltung, diesmal der *Impulsantwortmatrix* e^{At} mit dem Eingangsvektor $v(t)$. Es liegt nahe, zu vermuten, dass auch für Dgln n -ter Ordnung die Lösung durch Faltung einer entsprechend definierten Impulsantwort zu finden ist. Dies ist in der Tat der Fall; es gilt der

Satz 1.2.7 *Das betrachtete lineare System werde durch eine Dgl n -ter Ordnung mit zugehöriger Impulsantwort h beschrieben. Die durch das Eingangssignal $u(t)$ erzeugte Lösung $x(t)$ ist durch die Faltung*

$$x(t) = \int_0^t h(t-\tau) u(\tau) d\tau \quad (1.60)$$

gegeben.

Beweis: Der Beweis ergibt sich leicht durch die Operatorschreibweise und durch das in Satz 1.2.3, p. 19, bewiesene Superpositionsprinzip. Die Impulsantwort ist definiert durch $L[\delta(t)] = h(t)$. Damit gilt auch $L[\delta(t - \tau)] = h(t - \tau)$, wie man sich anhand der Substitution $\sigma = t - \tau$ leicht erzeugt. Da $u(t) = \int_{-\infty}^{\infty} u(\tau)\delta(t - \tau)d\tau$ erhält man

$$\begin{aligned} L[u(t)] &= L\left(\int_{-\infty}^{\infty} u(\tau)\delta(t - \tau)d\tau\right) \\ &= \int_{-\infty}^{\infty} u(\tau)L[\delta(t - \tau)]d\tau = \int_{-\infty}^{\infty} u(\tau)h(t - \tau)d\tau. \end{aligned}$$

□

Wählt man für u einen hinreichend kurzen Impuls, so ist die Antwort g eine Approximation der Impulsantwort selbst:

Satz 1.2.8 *Ein lineares System sei durch eine Impulsantwort $h(t)$ repräsentiert, und es sei*

$$\int h(t)dt = H(t) + c$$

c eine Integrationskonstante. Das Eingangssignal sei durch einen Rechteckpuls, d.h. durch

$$u(t) = \begin{cases} m/\Delta t, & t \in (t_0, t_0 + \Delta t) \\ 0, & t \notin (t_0, t_0 + \Delta t) \end{cases} \quad (1.61)$$

definiert. Dann gilt

$$\lim_{\Delta t \rightarrow 0} \int_{-\infty}^{\infty} h(t - \tau)u(\tau)d\tau = h(t) \quad (1.62)$$

Beweis: Es ist

$$\int_{-\infty}^{\infty} h(t - \tau)u(\tau)d\tau = \frac{1}{\Delta t} \int_{t_0}^{t_0 + \Delta t} h(t - \tau)d\tau, \quad t > \Delta t$$

Es genügt, den Fall $t > \Delta t$ zu betrachten, da ja die Reaktion auf einen Impuls von beliebig kurzer Dauer betrachtet werden soll. Substituiert man $\xi = t - \tau$, so hat man wegen $d\xi/dt = -1$

$$\int_{t_0}^{t_0 + \Delta t} h(t - \tau)d\tau = - \int_{t - t_0}^{t - t_0 - \Delta t} h(\xi)d\xi = \int_{t - t_0 - \Delta t}^{t - t_0} h(\xi)d\xi = H(t - t_0) - H(t - t_0 - \Delta t)$$

Aus der Definition von H folgt

$$\lim_{\Delta t \rightarrow 0} \frac{H(t + \Delta t) - H(t)}{\Delta t} = h(t).$$

Dann ist

$$\frac{1}{\Delta t} \int_{t_0}^{t_0 + \Delta t} h(t - \tau)d\tau = \frac{H(t - t_0) - H(t - t_0 - \Delta t)}{\Delta t} \rightarrow h(t - t_0), \quad \Delta t \rightarrow 0$$

Für $t_0 = 0$ (d.h. setzt man den Anfangszeitpunkt des Impulses gleich Null), so erhält man also gerade die Impulsantwort h für hinreichend kleinen Wert von Δt . Für $t > t_0 > 0$ erhält man ebenfalls die Impulsfunktion; man muß nur $\xi = t - t_0$ als neue Variable einführen.

□

Anmerkung: Es sei noch einmal daraufhingewiesen, dass die Intensität des Impulses proportional zu $1/\Delta t$ gegeben ist, also reziprok zur Dauer Δt gewählt wird. Denn nur dann ergibt sich für $\Delta t \rightarrow 0$ ein Dirac-Impuls, zu dessen definierenden Eigenschaften es ja gehört, dass das Integral über diese Funktion gleich 1 ist.

Nun wird in Experimenten der Wert von Δt nicht wirklich *beliebig* nahe bei Null liegen, sondern einen definitiv von Null verschiedenen Wert haben. Die Frage ist dann, wie gut für vorgegebenen Wert von Δt die Approximation für $h(t)$ ist. Es werde zunächst ein Beispiel betrachtet:

Beispiel 1.2.7 Rechteckpuls Für ein System erster Ordnung ist die Impulsantwort durch $h(t) = e^{at}$, $a \in \mathbb{R}$, gegeben. Es sei $u(t)$ insbesondere ein *Rechtecksignal* der Intensität $1/\Delta t$, wobei Δt die Dauer des Signals ist, vergl. (1.61). Die Antwort auf $u(t)$ ist dann

$$g(t) = \frac{1}{\Delta t} \int_{t_0}^{t_0+\Delta t} e^{a(t-\tau)} d\tau, \quad (1.63)$$

Es ist

$$\begin{aligned} \int_{t_0}^{t_0+\Delta t} e^{a(t-\tau)} d\tau &= e^{at} \int_{t_0}^{t_0+\Delta t} e^{-a\tau} d\tau \\ &= e^{at} \left(-\frac{1}{a} e^{-a\tau} \Big|_{t_0}^{t_0+\Delta t} \right) \\ &= \frac{e^{at}}{a} \left(e^{-at_0} - e^{-a(t_0+\Delta t)} \right) \\ &= \frac{1}{a} \left(e^{a(t-t_0)} - e^{a(t-t_0-\Delta t)} \right) \end{aligned}$$

Nun ist

$$e^{a(t-t_0-\Delta t)} = e^{a(t-t_0)} e^{-a\Delta t},$$

so dass

$$\int_{t_0}^{t_0+\Delta t} e^{a(t-\tau)} d\tau = \frac{1}{a} e^{a(t-t_0)} (1 - e^{-a\Delta t})$$

Für hinreichend kleinen Wert von Δt gilt aber

$$e^{-a\Delta t} \approx 1 - a\Delta t,$$

so dass

$$\begin{aligned} g(t) &= \frac{1}{\Delta t} \int_{t_0}^{t_0+\Delta t} e^{a(t-\tau)} d\tau \\ &= \frac{1}{a\Delta t} e^{a(t-t_0)} (1 - e^{-a\Delta t}) \approx \frac{1}{a\Delta t} e^{a(t-t_0)} (1 - 1 + a\Delta t) \\ &= \frac{1}{a\Delta t} e^{a(t-t_0)} a\Delta t = e^{a(t-t_0)}, \quad t > t_0 \end{aligned} \quad (1.64)$$

Für Werte von Δt , für die die Approximation $e^{a(t-t_0)} \approx 1 - a\Delta t$ hinreichend gut ist, approximiert auch die Reaktion auf den Impuls mit endlicher Ausdehnung Δt die Impulsantwort $h(t-t_0) = e^{a(t-t_0)}$. Dieser Wert von Δt wird von dem des Parameters a abhängen: Wegen $e^{-x} \approx 1 - x$ muß, mit $x = a\Delta t$, Δt umgekehrt proportional zu a sein.

□

Beispiel 1.2.8 (Schrittfunktion) Es sei jetzt $u(t)$ eine Schrittfunktion, so dass

$$u(t) = \begin{cases} m, & t > 0 \\ 0, & t \leq 0. \end{cases} \quad (1.65)$$

Dann ist die Antwort durch

$$g(t) = m \int_0^t h(t-\tau) d\tau = H(t) - H(0) \quad (1.66)$$

gegeben. Offenbar ist $dg(t)/dt = h(t)$. Bestimmt man also einerseits die Impulsfunktion, andererseits die Schrittfunktion empirisch, so kann man testen, ob sich anhand der (numerisch bestimmten) Ableitung der empirisch bestimmten Schrittfunktion die empirisch bestimmte Impulsfunktion "vorhersagen" läßt. Diese Vorhersage ist ein Test für die Güte der Messungen einerseits, aber andererseits auch der Hypothese, dass man es überhaupt mit einem linearen System zu tun hat. □

Die Frage ist zunächst, wie die Impulsantwort $h(t)$ für ein System n -ter Ordnung gefunden werden kann, wenn die Koeffizienten a_k und b_j gegeben sind. Da wir später insbesondere an der Analyse von Systemen mit unbekannter Struktur interessiert sind, müssen wir uns ebenfalls mit der Frage beschäftigen, wie man von einer empirisch bestimmten Impulsantwort auf die Systemstruktur zurückschließen kann. Die Begriffsbildungen und Betrachtungen im folgenden Abschnitt dienen u.a. der Vorbereitung der Beantwortung dieser Frage.

1.2.10 System- und Transferfunktionen

Allgemeine Resultate

Es soll zuerst an den Begriff der Fourier- und der Laplacetransformierten erinnert werden (s. Anhang). Ist $f(t)$ irgendeine reelle Funktion, die der Bedingung

$$\int_{-\infty}^{\infty} |f(t)| dt < \infty$$

genügt, so heißt

$$F(j\omega) = \int_{-\infty}^{\infty} f(t) e^{-j\omega t} dt, \quad j = \sqrt{-1} \quad (1.67)$$

Fouriertransformierte der Funktion f . $F(j\omega)$ ist eine komplexe Funktion, so dass $F(j\omega) = F_1(\omega) + jF_2(\omega)$ gilt, wobei F_1 und F_2 reelle Funktionen von ω sind. Der

Betrag von F ist dann

$$A(\omega) = |F(j\omega)| = \sqrt{F_1^2(\omega) + F_2^2(\omega)}, \quad (1.68)$$

und

$$\varphi(\omega) = \tan^{-1}[F_2(\omega)/F_1(\omega)]. \quad (1.69)$$

Dann läßt sich F in der Form

$$F(j\omega) = |F(j\omega)|e^{j\varphi(\omega)} = A(\omega)e^{j\varphi(\omega)} \quad (1.70)$$

darstellen. $A(\omega)$ ist die Amplitude, mit der die Frequenz $\omega = 2\pi f$ in $f(t)$ enthalten ist, und $\varphi(\omega)$ gibt die Phase für diese Frequenz an. Ist $F(j\omega)$ gegeben, so läßt sich daraus durch inverse Transformation die Funktion $f(t)$ bestimmen; es gilt

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(j\omega) e^{j\omega t} d\omega. \quad (1.71)$$

Die Laplacetransformierte von f ist durch

$$F(s) = L(f(t)) = \int_0^{\infty} f(t) e^{-st} dt \quad (1.72)$$

gegeben, wobei $s = \alpha + j\omega \in \mathbb{C}$, d.h. s sei eine komplexe Zahl. Die inverse Transformation ist dann

$$f(t) = \int_0^{\infty} F(s) e^{st} ds \quad (1.73)$$

Definition 1.2.5 *Es sei $h(t)$ die Impulsantwort eines linearen Systems mit konstanten Koeffizienten. Dann heißt die Fouriertransformierte H von h , also*

$$H(j\omega) = \int_0^{\infty} h(t) e^{-j\omega t} dt, \quad j = \sqrt{-1} \quad (1.74)$$

die Systemfunktion des Systems, und die Laplacetransformierte

$$F(s) = \int_0^{\infty} h(t) e^{-st} dt \quad (1.75)$$

heißt Transferfunktion des Systems.

Die Systemfunktion wird im Allgemeinen eine komplexe Funktion der Form

$$H(j\omega) = H_1(\omega) + jH_2(\omega) \quad (1.76)$$

sein, wobei H_1 der Real- und H_2 der Imaginärteil ist; H_1 und H_2 sind reelle Funktionen. Die Definition von Amplituden- und Phasenfunktionen (1.68) und (1.69) übertragen sich natürlich.

Satz 1.2.9 Gegeben sei ein lineares System mit konstanten Koeffizienten mit der Impulsantwort $h(t)$. Es seien $u(t)$ ein Eingangssignal mit der Laplacetransformierten $U(s)$, und $g(t)$ sei das zugehörige Ausgangssignal mit der Laplacetransformierten $G(s)$. $h(t)$ habe die Laplacetransformierte $H(s)$. Dann gilt

$$G(s) = H(s)U(s). \quad (1.77)$$

Für die Fouriertransformierten $G(j\omega)$, $H(j\omega)$ und $U(j\omega)$ gilt die analoge Beziehung

$$G(j\omega) = H(j\omega)U(j\omega). \quad (1.78)$$

Beweis: Die Laplacetransformierte von g ist

$$\begin{aligned} G(s) &= \int_0^\infty g(t)e^{-st}dt = \int_0^\infty \left(\int_{-\infty}^\infty u(\tau)h(t-\tau)d\tau \right) e^{-st}dt \\ &= \int_{-\infty}^\infty u(\tau) \int_0^\infty h(t-\tau)e^{-st}d\tau dt \\ &= \int_{-\infty}^\infty u(\tau) \int_0^\infty h(v)e^{-s(v+\tau)}dv, \quad v \stackrel{!}{=} t - \tau \\ &= \int_{-\infty}^\infty u(\tau)e^{-s\tau}d\tau \int_{-\infty}^\infty h(v)e^{-sv}dv \\ &= U(s)H(s) = H(s)U(s). \end{aligned}$$

Für die Fouriertransformierte $G(j\omega)$ läuft der Beweis analog. □

Aus (1.77) und (1.78) folgen sofort die Beziehungen

$$H(s) = \frac{G(s)}{U(s)} \quad (1.79)$$

und

$$H(j\omega) = \frac{G(j\omega)}{U(j\omega)} \quad (1.80)$$

Die Systemfunktion $H(j\omega)$ ist eine komplexe Funktion, d.h. es gibt zwei reelle Funktionen $P(\omega)$ und $Q(\omega)$ derart, dass

$$H(j\omega) = P(\omega) + jQ(\omega). \quad (1.81)$$

Dies folgt aus der Definition von H als Fouriertransformierter der Impulsantwort h .

Definition 1.2.6 Es sei $h(t)$ die Impulsantwort eines linearen Systems mit konstanten Koeffizienten. Die Laplacetransformierte $H(s)$ von $h(t)$ heißt Transferfunktion des Systems, die Fouriertransformierte $H(j\omega)$ heißt Systemfunktion oder Frequenzgang. $A(\omega) = |H(j\omega)|$ insbesondere heißt Amplitudengang und $\varphi(\omega)$ heißt Phasengang des Systems.

Definition 1.2.7 Es sei $h(t)$ die Impulsantwort eines linearen Systems mit konstanten Koeffizienten. Das System heißt kausal, wenn $h(t) = 0$ für alle $t < 0$.

Eine System ist also dann kausal, wenn das System nicht auf eine Störung reagieren kann, bevor diese Störung nicht auch tatsächlich eingewirkt hat. Solange t als Zeitvariable interpretiert wird, haftet dieser Begriffsbildung etwas Triviales an. Aber formal kann t auch durch eine Ortsvariable x ersetzt sein; in diesem Fall ist es oft durchaus sinnvoll, nichtkausale Systeme zu betrachten; Nichtkausalität bedeutet ja nur, dass $h(x) \neq 0$ für $x < 0$ sein darf.

Der Ausdruck *Frequenzgang* ergibt sich aus der Aussage des folgenden Satzes:

Satz 1.2.10 *L repräsentiere ein lineares, kausales System mit konstanten Koeffizienten. Das Eingangssignal sei $\sin(\omega t)$ für $-\infty < t < \infty$. Dann gilt*

$$g(t) = L(\sin(\omega t)) = |H(j\omega)| \sin(\omega t + \varphi(\omega)). \quad (1.82)$$

Beweis: Nach (1.60) ist die Antwort des Systems durch das Faltungsintegral

$$g(t) = \int_0^\infty h(t - \tau) \sin(\omega \tau) d\tau = \frac{1}{2j} \left(\int_0^\infty h(t - \tau) e^{j\omega \tau} d\tau - \int_0^\infty h(t - \tau) e^{-j\omega \tau} d\tau \right)$$

gegeben, wobei von der Eulerschen Relation $\sin \omega t = (e^{j\omega t} - e^{-j\omega t})/2j$ Gebrauch gemacht wurde. Mit der Substitution $\sigma = t - \tau$ erhält man dann

$$g(t) = \frac{1}{2j} \left(e^{j\omega t} \int_0^\infty h(\sigma) e^{-j\omega \sigma} d\sigma - e^{-j\omega t} \int_0^\infty h(\sigma) e^{j\omega \sigma} d\sigma \right). \quad (1.83)$$

Aber

$$\int_0^\infty h(\sigma) e^{-j\omega \sigma} d\sigma = H(j\omega) = |H(j\omega)| e^{j\varphi(\omega)}, \quad (1.84)$$

$$\int_0^\infty h(\sigma) e^{j\omega \sigma} d\sigma = H(-j\omega) = |H(j\omega)| e^{-j\varphi(\omega)} \quad (1.85)$$

Mithin ist

$$g(t) = \frac{1}{2j} \left(e^{j\omega t} |H(j\omega)| e^{j\varphi(\omega)} - e^{-j\omega t} |H(j\omega)| e^{-j\varphi(\omega)} \right),$$

woraus sich sofort (1.82) ergibt. □

Anmerkung: Es ist nach den Eulerschen Relationen und dem Superpositionsprinzip

$$L(\sin \omega t) = \frac{1}{2j} L(e^{j\omega t}) - \frac{1}{2j} L(e^{-j\omega t})$$

Aus (1.83), (1.84) und (1.85) folgt dann

$$L(e^{j\omega t}) = H(j\omega) e^{j\omega t}, \quad L(e^{-j\omega t}) = H(j\omega) e^{-j\omega t} \quad (1.86)$$

Nach (1.82) ist die Antwort auf eine Sinusschwingung der Frequenz ω wieder eine Sinusschwingung mit der gleichen Frequenz, allerdings um den Betrag $\varphi(\omega)$

phasenverschoben. Die Eingangsamplitude ist 1, die Ausgangsamplitude $A(\omega) = |H(j\omega)|$ ist nicht notwendig ebenfalls gleich 1, sie wird im allgemeinen ebenso wie φ von ω abhängen. Ist die Eingangsamplitude gleich c , so ist die Ausgangsamplitude gleich $mA(\omega)$. Für $A(\omega) > 1$ verstärkt das System die Amplitude (Gain), für $A(\omega) < 1$ wird sie abgeschwächt (Attenuation). $H(j\omega)$ liefert also für jedes ω die zugehörige Ausgangsamplitude $A(\omega)$ und die Phase $\varphi(\omega)$; daher der Ausdruck *Frequenzgang*.

Es soll noch einmal darauf hingewiesen werden, dass in Satz 1.2.10 angenommen wurde, dass $\sin(\omega t)$ auf dem gesamten Intervall $(-\infty, t]$ oder $[0, \infty)$ auf das System einwirkt; (1.82) beschreibt dann die Reaktion im *eingeschwungenen Zustand*. Wird $u(t)$ in $t = 0$ eingeschaltet und betrachtet man die Reaktion in $\infty > t > 0$, so werden außer $A(\omega) \sin(\omega t + \varphi)$ noch *transiente Effekte* existieren, die aber für $t \rightarrow \infty$ im allgemeinen gegen Null gehen. Auf die transienten Effekte wird weiter unten eingegangen. Zunächst soll (1.82) benutzt werden, um einen ersten Ausdruck für $H(j\omega)$ zu erhalten. Es gilt der

Satz 1.2.11 *Das durch L repräsentierte lineare System sei durch die Dgl n -ter Ordnung (1.54) (p. 28), d.h. durch*

$$\sum_{k=0}^n a_k D^k z(t) = \sum_{j=1}^m b_j D^j u(t)$$

charakterisiert. Dann gilt

$$H(j\omega) = \frac{b_m(j\omega)^m + b_{m-1}(j\omega)^{m-1} + \dots + b_0(j\omega)}{a_n(j\omega)^n + a_{n-1}(j\omega)^{n-1} + \dots + a_0(j\omega)}. \quad (1.87)$$

Beweis: Man betrachte das Eingangssignal $u(t) = \sin \omega t = (e^{j\omega t} - e^{-j\omega t})/2j$. Wegen des Superpositionsprinzips kann man die Reaktionen $g_1(t)$ und $g_2(t)$ des Systems auf $u_1(t) = e^{j\omega t}$ und $u_2(t) = e^{-j\omega t}$ gesondert betrachten und dann die Antwort auf $u(t)$ gemäß $g(t) = (g_1(t) + g_2(t))/2j$ bestimmen. Setzt man $u_1(t)$ in die Differentialgleichung ein, so erhält man wegen $du_1(t)/dt = j\omega e^{j\omega t}$, $d^2 u_1(t)/dt^2 = (j\omega)^2 e^{j\omega t}$ etc. Aus (1.86) weiß man aber, dass $L[e^{j\omega t}] = g_1(t) = H(j\omega)e^{j\omega t}$. Also folgt

$$\sum_{k=0}^n a_k (j\omega)^k H(j\omega) = \sum_{j=0}^m b_j (j\omega)^j,$$

und dieser Ausdruck führt sofort auf (1.87). Für $u_2(t) = e^{-j\omega t}$ wird man auf das gleiche Resultat geführt. $\square$

Gemäß (1.58) ist die Systemfunktion eine *rationale Funktion*, d.h. der Quotient zweier Polynome endlicher Ordnung. Betrachtet man statt des Eingangssignals $u_1(t) = e^{j\omega t}$ das Signal $u_1(t) = e^{st}$ mit $s = \alpha + j\omega$, so erhält man auf die gleiche Weise den Ausdruck

$$H(s) = \frac{b_m s^m + \dots + b_0 s}{a_n s^n + \dots + a_0 s} \quad (1.88)$$

für die Transferfunktion, die sich also gleichermaßen als eine rationale Funktion erweist. Es werde noch das folgende Korollar angemerkt:

Korollar: Die Transferfunktion eines Systems sei eine rationale Funktion, d.h. $H(s)$ sei durch (1.88) gegeben. Dann wird das System durch eine Dgl n -ter Ordnung der Form (1.54), d.h. durch

$$\sum_{k=0}^n a_k D^k z(t) = \sum_{j=1}^m b_j D^j u(t)$$

beschrieben.

Beweis: Es folgt aus (1.88), dass

$$(a_n s^n + \dots + a_0 s)G(s) = (b_m s^m + \dots + b_0 s)U(s).$$

Man bildet jetzt die inverse Transformierte auf beiden Seiten und erhält

$$a_m g^{(n)} + \dots + a_0 g(t) = b_m u^{(m)}(t) + \dots + b_0 u(t).$$

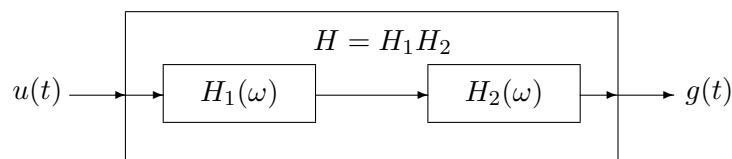
□

Entscheidet man sich also, eine *empirisch* bestimmte Transferfunktion als rationale Funktion zu interpretieren, so hat man damit auch die Form der beschreibenden Dgl bestimmt: es muß sich um eine Dgl n -ter Ordnung handeln. dass dieses wiederum als ein System von Dgln erster Ordnung geschrieben werden kann, wurde in Abschnitt 2.2 gezeigt.

Bevor hergeleitet werden kann, wie der funktionale Bauplan des Systems bestimmt werden kann, sollen die Grundprinzipien der Zusammenschaltung von Systemen angegeben werden. Denn die Komponenten eines Systems können ja selbst als System aufgefaßt werden, und damit sind die Komponenten ebenfalls durch Impulsantworten bzw. Transferfunktionen charakterisiert. Die Analyse besteht dann in der Anwendung der folgenden Prinzipien.

1. **Die Hintereinander- oder Reihenschaltung von Systemen:** Das Ausgangssignal eines Systems wird dann zum Eingangssignal des folgenden Systems. Ist $g_1(t)$ das Ausgangssignal des ersten Systems und ist $h_2(t)$ die Im-

Abbildung 1.3: Reihenschaltung



pulsantwort des zweiten Systems, so ist $g_2(t)$, die Antwort des zweiten Systems,

durch die Faltung von $g_1(t)$ und $h_2(t)$ gegeben. Dieser Faltung entspricht aber das Produkt der Laplace-Transformationen von x_1 und h_2 (bzw. der Fouriertransformationen dieser Funktionen). Also gilt

$$G_2(s) = G_1(s)H_2(s). \quad (1.89)$$

Will man also die Beziehung zwischen dem Eingangssignal $u(t)$ (für das erste System) mit der Laplace-Transformierten $U(s)$ und dem Ausgangssignal des zweiten Systems mit der Laplace-Transformierten $G_2(s)$, so erhält man aus (1.89) wegen $G_1(s) = U(s)H_1(s)$

$$G_2(s) = H_1(s)H_2(s)U(s). \quad (1.90)$$

Für n rückwirkungsfrei hintereinandergeschaltete Systeme hat man dann sofort

$$G_n(s) = \left(\prod_{i=1}^n H_i(s) \right) U(s) = H(s)U(s) \quad (1.91)$$

mit

$$H(s) = \left(\prod_{i=1}^n H_i(s) \right) \quad (1.92)$$

Die hintereinandergeschalteten Systeme können alle als kausal vorausgesetzt werden, so dass $h_i(t) = 0$ für $t \leq 0$ gilt. Unter sehr allgemeinen Nebenbedingungen gilt dann

$$H(j\omega) = \prod_{i=1}^n H_i(j\omega) = A e^{-j\mu\omega - \omega^2\sigma^2/2}, \quad j = \sqrt{-1} \quad (1.93)$$

Dies ist die Fouriertransformierte der Gauß-Funktion; die Impulsantwort des Gesamtsystems für hinreichend großes n ist dann also durch

$$h(t) \approx \frac{A}{\sigma\sqrt{2\pi}} e^{-(t-\mu)^2/2\sigma^2} \quad (1.94)$$

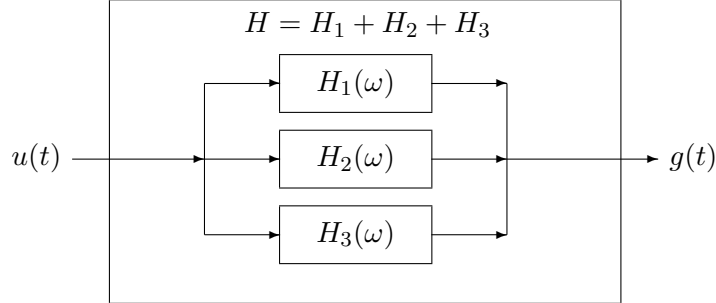
gegeben. Diese Aussage entspricht dem zentralen Grenzwertsatz in der Statistik. (vergl. Papoulis (1968), p. 779). $\square$

- 2. Parallelschaltung von Systemen:** Es werde nun die Situation in Abb. 1.4 betrachtet. Wieder sei $u(t)$ das Eingangssignal, diesmal aber für beide Systeme L_1 und L_2 mit den Transferfunktionen H_1 und H_2 , und die Ausgänge der beiden Systeme seien $g_1(t)$ und $g_2(t)$. Das Ausgangssignal $g(t)$ sei durch $g(t) = g_1(t) + g_2(t)$ definiert.

Es sei wieder $U(s)$ die Laplace-Transformierte von $u(t)$, H_1 und H_2 seien die Transferfunktionen von L_1 bzw. L_2 , und $G_1(s)$, $G_2(s)$ seien die Laplacetransformierten von $g_1(t)$ und $g_2(t)$. Für die Transferfunktion $H(s)$ des Gesamtsystems gilt

$$H(s) = H_1(s) + H_2(s). \quad (1.95)$$

Abbildung 1.4: Parallelschaltung



Denn aus $g(t) = g_1(t) + g_2(t)$ folgt durch Laplace-Transformation

$$\begin{aligned} L(g(t)) = G(s) &= L(g_1(t)) + L(g_2(t)) = G_1(s) + G_2(s) \\ &= U(s)H_1(s) + U(s)H_2(s), \end{aligned}$$

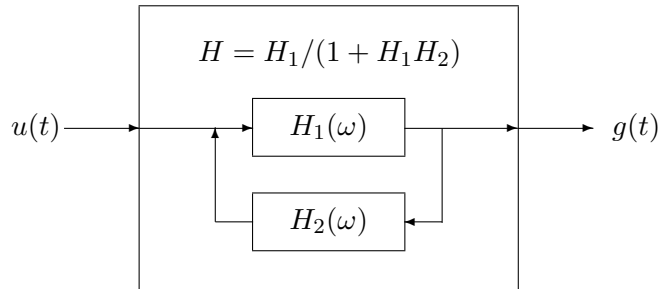
d.h.

$$G(\omega) = U(\omega)[H_1(\omega) + H_2(\omega)],$$

woraus sofort die Behauptung folgt.

3. **Feedback-Systeme:** Aus (1.89) und (1.95) lässt sich sofort die Transferfunktion eines Feedback-Systems herleiten, vergl. Abb. 1.5. Die Systeme L_1 und L_2

Abbildung 1.5: Feedback-Schaltung



mögen wieder die Transferfunktionen H_1 und H_2 haben. Gesucht ist die Transferfunktion H , über die y und u miteinander verknüpft sind. $u_1(t)$ geht aus $u(t)$ durch Subtraktion von $u_3(t)$ hervor, d.g. $u_1(t) = u(t) - u_3(t)$. $u_3(t)$ wiederum ist der Ausgang des Systems L_2 mit dem Eingang $u_2(t)$. Es sei $U(s)$ die L -Transformierte von $u(t)$, und $U_j(s)$ seien die L -Transformierten von $u_j(t)$. Dann folgt $U_1(s) = U(s) - U_3(s)$, $U_2(s) = H_1(s)U_1(s)$, $U_3(s) = H_2(s)U_2(s)$. Gleichzeitig gilt $G(s) = U_2(s)$. Wir suchen die Beziehung zwischen $U_2(s)$ und $U(s)$. Nach den eben angegebenen Beziehungen gilt

$$U_2(s) = H_1(s)U_1(s) = H_1(s)(U(s) - U_3(s)) = H_1(s)(U(s) - H_2(s)U_2(s)),$$

woraus man sofort

$$G(s) = U_2(s) = U(s) \frac{H_1(s)}{1 + H_1(s)H_2(s)} \quad (1.96)$$

gewinnt, d.h. die Transferfunktion eines einfachen Feedback-Systems ist durch

$$H(s) = \frac{H_1(s)}{1 + H_1(s)H_2(s)} \quad (1.97)$$

gegeben. Da H_1 und H_2 selbst wieder Feedback-Systeme repräsentieren können, kann man rekursiv durch wiederholte Anwendung von (1.97) $H(s)$ in Termen der einzelnen konstituierenden Systeme anschreiben. Generell gilt (1.60), p. 31; andererseits gilt (1.92): mit $H_3(s) = 1/(1 - H_1(s)H_2(s))$ hat man ja $H(s) = H_1(s)H_3(s)$, also eine Produktdarstellung von $H(s)$. Unter Berücksichtigung von (1.95) gilt also, dass eine als rationale Funktion repräsentierte Transferfunktion stets kann als Summe von Transferfunktionen bzw. von Produkten von Transferfunktionen dargestellt werden kann. (1.60), p. 31, ist also ebenfalls in dieser Form darstellbar. Die Darstellung erlaubt es, von einer Transferfunktion auf die Zusammenschaltung von Subsystemen zu schließen. Um diesen Schluß konkret durchführen zu können, macht man von der im nächsten Abschnitt kurz dargestellten Partialbruchzerlegung Gebrauch.

Beispiel 1.2.9 Es werde ein einfaches Beispiel für eine Impulsfunktion, die sich aus Impulsfunktionen für Teilsysteme zusammensetzt, betrachtet. Nach Gleichung (1.44) läßt sich jede Komponente eines Vektors von Impulsantworten und damit eine Impulsantwort selbst als Summe von Termen der Art $a \exp(\lambda t)$ schreiben, wobei a und λ im allgemeinen komplexe Zahlen sind. Dies bedeutet, dass die Impulsantwort einer Parallelschaltung solcher "elementarer" Systeme entspricht. Hintereinanderschaltungen entsprechen Faltungen von Impulsantworten. Hier werde die Impulsantwort

$$h(t) = a_1 e^{\lambda_1 t} + a_2 e^{\lambda_2 t} - b_1 e^{\mu_1 t} - b_2 e^{\mu_2 t} \quad (1.98)$$

mit — betrachtet. Die Parameter sind bewußt so gewählt worden, dass die resultierende Funktion nicht "schön" aussieht, damit der Effekt der einzelnen Komponenten deutlich wird. Die Funktion wird in Abb. zusammen mit der zugehörigen Schrittantwort dargestellt. □

1.2.11 Minimum-Phasen-Systeme

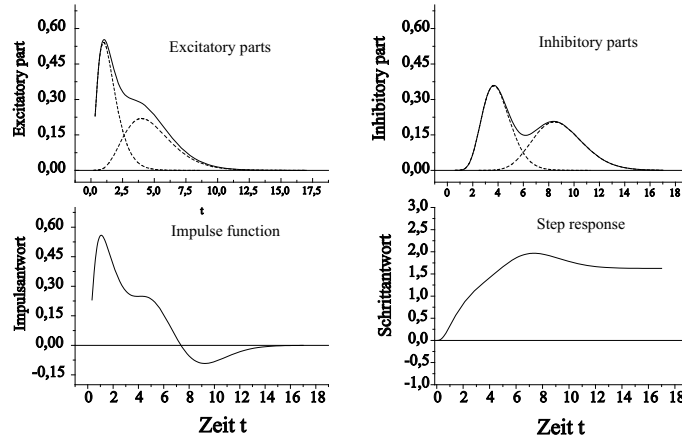
Bekanntlich heißt eine Funktion f *gerade*, wenn $f(-t) = f(t)$ gilt, und *ungerade*, wenn $f(-t) = -f(t)$ gilt. Es gilt der

Satz 1.2.12 Für jede Funktion f lassen sich eine ungerade Funktion f_u und eine gerade Funktion f_g finden derart, dass

$$f(t) = f_u(t) + f_g(t), \quad (1.99)$$

d.h. die Funktion f läßt sich als Summe einer geraden und einer ungeraden Funktion darstellen.

Abbildung 1.6: Impuls- und Schrittantwort



Beweis: Man kann stets $f(t) = f(t)/2 + f(-t)/2 + f(t)/2 - f(-t)/2$ schreiben. Etwas umgeordnet ergibt sich

$$f(t) = \frac{1}{2}(f(t) + f(-t)) + \frac{1}{2}(f(t) - f(-t)).$$

Definiert man

$$\frac{1}{2}(f(t) + f(-t)) \stackrel{\text{def}}{=} f_g(t), \quad (1.100)$$

$$\frac{1}{2}(f(t) - f(-t)) \stackrel{\text{def}}{=} f_u(t), \quad (1.101)$$

so erhält man gerade $f(t) = f_u(t) + f_g(t)$. dass f_g und f_u tatsächlich gerade bzw. ungerade Funktionen sind, ergibt sich leicht durch Einsetzen von $-t$ für t . $\square$

Anmerkungen:

1. Mit $\text{sgn } t$ wird i.a. das Vorzeichen von t bezeichnet. In bezug auf (1.100) und (1.101) findet man dann z.B. $f_u(t) = \text{sgn } t f_g(t)$.
2. Die Fouriertransformierte F von f ist dann die Summe der Fouriertransformierten F_u und F_g von f_u und f_g , denn

$$\begin{aligned} F(\omega) &= \int_{-\infty}^{\infty} f(t) e^{-j\omega t} dt \\ &= \int_{-\infty}^{\infty} f_u(t) e^{-j\omega t} dt + \int_{-\infty}^{\infty} f_g(t) e^{-j\omega t} dt \end{aligned} \quad (1.102)$$

Definition 1.2.8 Es sei f eine reelle Funktion. Dann heißt

$$\hat{f}(t) = \frac{1}{\pi} P \int_{-\infty}^{\infty} \frac{f(\tau)}{t - \tau} d\tau \quad (1.103)$$

die Hilbert-Transformierte von f ; P ist der Hauptwert des Integrals.

Anmerkung: Der *Hauptwert* von $g(\tau) = f(t)/(t - \tau)$ ist

$$P \int_{-\infty}^{\infty} g(\tau) d\tau = \lim_{\epsilon \rightarrow 0} \left(\int_{-\infty}^{-\epsilon} g(\tau) d\tau + \int_{\epsilon}^{\infty} g(\tau) d\tau \right)$$

Es gilt der

Satz 1.2.13 *Es sei $h(t)$ die Impulsantwort eines kausalen Systems. Dann bilden der Realteil $P(\omega)$ und der Imaginärteil $Q(\omega)$ der zugehörigen Systemfunktion $H(j\omega)$ ein Paar von Hilberttransformierten, d.h.*

$$P(\omega) = \frac{1}{\pi} \int_{-\infty}^{\infty} \frac{Q(\eta)}{\eta - \omega} d\eta, \quad Q(\omega) = \frac{1}{\pi} \int_{-\infty}^{\infty} \frac{P(\eta)}{\eta - \omega} d\eta. \quad (1.104)$$

Beweis: Nach Satz 1.2.12 kann die Impulsantwort als Summe einer geraden und einer ungeraden Funktion angeschrieben werden, so dass $h(t) = h_g(t) + h_u(t)$, wobei wieder $h_g(t) = \frac{1}{2}[h(t) + h(-t)]$, $h_u(t) = \frac{1}{2}[h(t) - h(-t)]$ gesetzt werden kann. Offenbar ist $h_u(t) = \text{sgn } t \, h_g(t)$, und so ist $h(t) = h_g(t)[1 + \text{sgn } t] = h_g(t) + h_g(t) \text{sgn } t$. Nach (1.102) ist dann die Fouriertransformierte $H(\omega)$ von $h(t)$ durch die Fouriertransformierte $H_g(\omega)$ von h_g plus der Faltung von H_g mit der Fouriertransformierten von $\text{sgn } t$ gegeben. Es sei $H_g(\omega) = F(h_g)$, und $F(h_g \text{sgn } t) = -H_g * j/\pi\omega$, also

$$H(\omega) = H_g(\omega) - \frac{j}{\pi\omega} * H_g(\omega).$$

Aber $-H_g(\omega) * j/\pi\omega$ ist gerade die Hilberttransformierte von $H_g(\omega)$. Allgemein gilt $H(\omega) = P(\omega) + jQ(\omega)$, so dass

$$P(\omega) + jQ(\omega) = H_g(\omega) - \frac{j}{\pi\omega} * H_g(\omega)$$

gelten muß. Dann ist $P(\omega) = H_g(\omega)$ und $Q(\omega) = -(j/\pi\omega) * H_g(\omega)$, und somit bilden P und Q ein Paar von Hilberttransformierten. $\square$

Die Aussage des Satzes 1.2.13 ist, dass der Real- und der Imaginärteil kausaler Systeme nicht unabhängig voneinander sind; die Abhängigkeit zwischen diesen beiden Komponenten der Systemfunktion läßt sich bei ihrer empirischen Bestimmung ausgenutzen.

Es sei nun $H(j\omega)$ die Systemfunktion eines kausalen Systems. H kann in der Form

$$H(j\omega) = P(\omega) + jQ(\omega) = A(\omega)e^{j\phi(\omega)} \quad (1.105)$$

dargestellt werden, mit

$$A(\omega) = |H(j\omega)| = \sqrt{P^2(\omega) + Q^2(\omega)}, \quad \phi(\omega) = \tan^{-1}(Q(\omega)/P(\omega))$$

Dann ist

$$\log H(j\omega) = \log A(\omega) + j\phi(\omega). \quad (1.106)$$

Dieser Ausdruck hat einen Aufbau, der analog zu dem von in (1.105) für $H(j\omega)$ gegebenen ist, denn $\log H(j\omega)$ ist ebenfalls eine komplexe Zahl, und $\log A(\omega)$ und $\phi(\omega)$ sind reell. Zwischen P und Q existieren aber die Beziehungen (1.104), d.h. P und Q sind ein Paar von Hilberttransformierten. Die Frage ist dementsprechend, ob nicht auch zwischen $\log A(\omega)$ und $\phi(\omega)$ diese Beziehung besteht. Dann müßte

$$\log A(\omega) = -\frac{1}{\pi} \int_{-\infty}^{\infty} \frac{\phi(u)}{u - \omega} du, \quad (1.107)$$

$$\phi(\omega) = \frac{1}{\pi} \int_{-\infty}^{\infty} \frac{\log A(u)}{u - \omega} du. \quad (1.108)$$

gelten (Solodovnikov (1960), p. 45). Der Vorteil der Existenz der Beziehungen (1.107) und (1.108) liegt auf der Hand: ist z.B. $A(\omega) = |H(j\omega)|$ bekannt, so läßt sich über (1.108) die Phasencharakteristik $\phi(\omega)$ berechnen.

Um die Bedingungen zu charakterisieren, unter denen die Gleichungen (1.107) und (1.108) gelten, muß man sich überlegen, dass eine Nullstelle von $H(j\omega)$ eine Singularität für $\log H(j\omega)$ bedeutet: $\log H(j\omega)$ wird dann $-\infty$, und damit auch $\log A(\omega)$, so dass (1.108) nicht mehr anwendbar wird. Es läßt sich zeigen, dass (1.107) und (1.108) nur gelten, wenn $H(j\omega)$ weder Pole noch Nullstellen in der unteren Halbebene aufweist.

Definition 1.2.9 *Lineare, kausale Systeme mit konstanten Koeffizienten, für die Beziehungen (1.107) und (1.108) gelten, heißen Minimum-Phasen-Systeme.*

Nach Satz 1.2.9 gilt $H(s) = G(s)/U(s)$ bzw. $H(j\omega) = G(j\omega)/U(j\omega)$, d.h. die Transfer- bzw. die Systemfunktion ist als Quotient zweier komplexer Funktionen, insbesondere zweier Polynome darstellbar (vergl. (1.58) und (1.59)). Aus der Algebra ist bekannt, dass sich ein Polynom als Produkt von komplexen Zahlen darstellen läßt. Es ist etwa

$$b_m s^m + \dots + b_0 s = \prod_{i=1}^m (s - \bar{s}_i), \quad a_n s^n + \dots + a_0 s = \prod_{j=1}^n (s - s_j)$$

wobei $\bar{s}_i$ und s_i die *Wurzeln* der entsprechenden Polynome sind. Aber die $s - \bar{s}_i$ und $s - s_i$ sind komplexe Zahlen, die sich in der Form $z = |z|e^{j\varphi}$ darstellen lassen, und ihr Produkt ist gleich dem Produkt ihrer Absolutbeträge (Moduli) sowie der Terme $e^{j\varphi}$. Der Phasenwinkel des Produkts ist gleich der Summe der Phasenwinkel, bei einem Quotienten gleich der Differenz der Phasenwinkel. Aber $G(s)$ ist schließlich der Quotient zweier komplexen Zahlen, und der korrespondierende Phasenwinkel ist gleich der Differenz der Phasenwinkel von Zähler- und Nennerpolynom. Für Minimalphasensysteme nimmt diese Differenz für alle ω den jeweils kleinstmöglichen Wert an. Eine ausführliche Diskussion findet man in Varjù (1977), p.123.

1.2.12 Transferfunktionen und Partialbruchzerlegung

Betrachtet man ein durch eine Dgl n -ter Ordnung beschriebenes System, so weiß man (vergl. (1.88)), dass die zugehörige Transferfunktion $H(s)$ die Form

$$H(s) = \frac{b_m s^m + b_{m-1} s^{m-1} + \dots + b_0 s}{a_n s^n + a_{n-1} s^{n-1} + \dots + a_0 s} \quad (1.109)$$

hat, wobei $s = \alpha + j\omega$, und die Systemfunktion hat die Form

$$H(j\omega) = \frac{b_m (j\omega)^m + b_{m-1} (j\omega)^{m-1} + \dots + b_0 (j\omega)}{a_n (j\omega)^n + a_{n-1} (j\omega)^{n-1} + \dots + a_0 (j\omega)}, \quad (1.110)$$

Damit ist $H(s)$ der Quotient zweier Polynome $P(s)$ und $Q(s)$. Wir beschränken uns im Folgenden auf $H(s)$; die Betrachtungen bezüglich $H(j\omega)$ sind analog. Um die inverse Transformation von $H(s)$, d.h. um die Impulsantwort zu finden, kann man nun $H(s) = P(s)/Q(s)$ in eine Summe von Teilbrüchen zerlegen, d.h. man führt eine *Partialbruchzerlegung* durch.

Nach Satz 1.2.11, Gleichung (1.58), p. 38, ist die Transfer- bzw. Systemfunktion der Quotient zweier Polynome in $s = \alpha + j\omega$. Es sei insbesondere $P(s) = a_n s^n + a_{n-1} s^{n-1} + \dots + a_0$, $Q(s) = b_m s^m + b_{m-1} s^{m-1} + \dots + b_0$. Aus der Algebra ist bekannt, dass ein Polynom in der Form eines Produktes geschrieben werden kann; man hat

$$P(s) = \prod_{i=1}^n (s - s_i), \quad Q(s) = \prod_{j=1}^m (s - s_j), \quad (1.111)$$

wobei s_i und s_j die *Nullstellen* von $P(s)$ und $Q(s)$ sind, d.h. $P(s_i) = 0$, $Q(s_j) = 0$. Nun ist

$$H(s) = \frac{P(s)}{Q(s)};$$

nimmt s den Wert einer Nullstelle von $P(s)$ an, so wird $H(s) = 0$; die s_i heißen dementsprechend die *Nullstellen* von $H(s)$. Nimmt s den Wert einer Nullstelle s_j von $Q(s)$ an, so wird $H(s)$ gleich unendlich. Die Nullstellen von $Q(s)$ heißen dann die *Pole* von $H(s)$.

Die Diskussion der Eigenschaften eines Systems ergibt sich aus den Nullstellen und Polen von H . Die Pole ergeben sich aus der Struktur von $Q(s)$, und dieses Polynom ergibt sich, wenn man die *homogenen Dgl*

$$Q(s) = a_n s^n + a_{n-1} s^{n-1} + \dots + a_0 s = 0$$

betrachtet. Man macht für diese Gleichung den Ansatz $x(t) = e^{st}$. Dann gilt $d^k x(t)/dt^k := x^{(k)} = s^k e^{st}$ und man erhält

$$\sum_{k=0}^n a_k x^{(k)}(t) x z e^{st} = e^{st} \sum_{k=0}^n a_k s^k = 0,$$

und dies kann nur gelten, wenn

$$Q(s) = \sum_{k=0}^n a_k s^k = 0$$

$Q(s)$ heißt auch das *charakteristische Polynom* der Dgl (1.54), p. 28. Die Bedingung $Q(s) = 0$ liefert dann gerade diejenigen Werte für s , für die e^{st} eine Lösung der homogenen Dgl ist. Die Nullstellen von $P(s)$ beziehen sich auf die Summe $\sum_{j=1}^n b_j u^{(j)}(t)$; setzt man hier $u(t) = e^{st}$, so ist $u^{(j)}(t) = s^j e^{st}$ und man hat

$$\sum_{j=0}^m b_j u^{(j)}(t) = \sum_{j=0}^m b_j s^j e^{st} = P(s) e^{st}.$$

Die Nullstellen von $P(s)$ sind dann diejenigen Werte von s , für die $u(t)$ identisch verschwindet.

Nun ist $Q(s)$ ein Polynom n -ten Grades, und deshalb gibt es gerade n mögliche Nullstellen. Dabei kann es sein, dass einige Nullstellen mehrfach, etwa ν_i -fach, vorkommen. Sie haben dann die *Mehrfachheit* ν_i . Gibt es r verschiedene Nullstellen mit den jeweiligen Mehrfachheiten $\nu_r \geq 1$, so muß insgesamt $\nu_1 + \dots + \nu_r = n$ gelten. Weiter gilt, dass eine Nullstelle entweder reell oder komplex ist. Ist eine Nullstelle s_j komplex, so existiert stets eine weitere komplexe Nullstelle s_k , die zu s_j konjugiert komplex ist, d.h. es gilt $s_j = a_j + j\omega_j$ und $s_k = \bar{s}_j = a_j - j\omega_j$. Wie weiter unten noch verdeutlicht werden wird, ist es gerade dieser Sachverhalt, der für die komplexen Nullstellen eine reelle Deutung bezüglich des Verhaltens des Systems impliziert.

Um die Eigenschaften des Systems aus den Nullstellen und Polen von $H(s)$ herzuleiten, wendet man die aus der Algebra bekannte *Partialbruchzerlegung* an: sie erlaubt es, den Quotienten zweier Polynome als Summe bestimmter Brüche anzuschreiben. Diese Summanden repräsentieren dann, nach den Ergebnissen von Abschnitt 1.2.12 über die Zusammenschaltung von Systemen, bestimmte parallelgeschaltete Untersysteme, deren Transferfunktion dann aus der Form der Summanden geschlossen werden kann. Wir beginnen mit dem einfachsten Fall:

a. Die Pole von $H(s)$ sind paarweise verschieden

Wie oben bereits angemerkt, können die Nullstellen von Polynomen reell oder komplex sein. In jedem Fall gilt nun

$$H(s) = \frac{P(s)}{Q(s)} = \frac{A_1}{s - s_1} + \frac{A_2}{s - s_2} + \dots + \frac{A_n}{s - s_n}. \quad (1.112)$$

Dabei sind die A_k , $k = 1, \dots, n$ Konstanten, die sich aus den a_i und den b_j ergeben. Da die a_i und b_j im Falle unbekannt sind, wenn man ein vorgegebenes System analysiert, soll auf die Berechnung der A_k hier nicht weiter eingegangen werden.

Aus (1.112) ergibt sich durch inverse Laplacetransformation sofort ein Ausdruck für die Impulsantwort $h(t)$. Dazu betrachten wir die inverse Transformation für den Ausdruck $A_i/(s - s_i)$; es ist

$$L^{-1} \left(\frac{A_i}{s - s_i} \right) = A_i e^{s_i t}, \quad (1.113)$$

und somit ist

$$h(t) = A_1 e^{s_1 t} + \dots + A_n e^{s_n t}. \quad (1.114)$$

Setzen wir $h_i(t) = A_i e^{s_i t}$, so gilt

$$h(t) = \sum_{i=1}^n h_i(t) = \sum_{i=1}^n A_i e^{s_i t}$$

Vergleicht man nun (1.113) (und damit (1.114)) mit (1.65), p. 34, so sieht an, dass das System als eine Parallelschaltung von (Teil-) Systemen aufgefaßt werden kann, die jeweils die Transferfunktion $H_i(s) = A_i/(s - s_i)$ und damit die Impulsantwort $A_i e^{s_i t}$ haben. Nun wurde in Abschnitt 2.2 gezeigt, dass ein System mit dieser Impulsantwort ein System erster Ordnung ist. Für das i -te Teilsystem muß also

$$\dot{z}_i - s_i z_i = u.$$

gelten. Nach Satz 1.2.2, p. 15, hat diese Dgl die Lösung

$$z_i(t) = e^{-s_i t} \int e^{s_i \tau} u(\tau) d\tau = e^{-s_i t}$$

für $u(t) = \delta(t)$ (der Faktor A_i ist hier vernachlässigt worden). Für das System kann dann das System von Dgln erster Ordnung angeschrieben werden:

$$\dot{z} = Az + bu, \quad (1.115)$$

wobei $A = \text{diag}(s_1, \dots, s_n)$, $b = (1, 1, \dots, 1)'$ und $x(t) = (a_1, \dots, a_n)z + du$ geschrieben werden. Die Tatsache, dass A eine Diagonalmatrix ist, bedeutet, dass die Variablen z_i *entkoppelt* sind, denn die $\dot{z}_i$ werden nur durch die z_i selbst, nicht aber durch z_j , $j \neq i$ beeinflusst. Die Variablen z_i werden gelegentlich auch als *kanonische Variablen* bezeichnet.

Zur weiteren Interpretation muß man nun zwischen zwei Fällen unterscheiden: (i) die s_i sind alle reell, (ii) es existieren auch komplexe Nullstellen.

1.2.13 Der Fall einfacher Pole

Alle Nullstellen sind reell Den einzelnen Summanden $h_i(t) = A_i e^{s_i t}$ können als Impulsantworten von Systemen erster Ordnung angesehen werden. Da die h_i reelle Funktionen sind, folgt, dass alle A_i ebenfalls reell sind. Der Transferfunktion $H(s)$ entspricht damit eine Parallelverschaltung von n Systemen erster Ordnung. Offenbar strebt $h(t)$ nur dann gegen Null, wenn die s_i alle negativ sind; in diesem Fall ist das System *stabil*. Ist auch nur eine der Nullstellen positiv, so bedeutet dies, dass der entsprechende Term auch bei geringster Anregung des Systems gegen unendlich strebt; das System ist dann *instabil*.

Es existieren komplexe Nullstellen Die Nullstelle s_i des Polynoms $Q(s)$, d.h. der Pol s_i von $H(s)$, sei nun komplex; es gelte $s_i = \alpha_i + j\omega_i$. Dann existiert ein weiterer Pol s_j derart, dass $s_j = \bar{s}_i = \alpha_i - j\omega_i$. Generell gelte wieder (1.112), allerdings sind jetzt die Koeffizienten A_i und A_j ebenfalls komplex. Es gilt insbesondere $A_j = \bar{A}_i$, d.h. A_j ist zu A_i konjugiert komplex, so dass $A_i = |A_i|e^{i\varphi_i}$, $A_j = |A_i|e^{-i\varphi_i}$. In (1.112) kann man nun die Terme mit komplexen und dazu konjugiert komplexen s_i ,

$\bar{s}_i$ zusammenfassen und die jeweiligen Summen $h_i(t) = A_i e^{s_i t} + A_j e^{s_j t}$ betrachten. Es gilt dann

$$h_i(t) = |A_i| e^{\alpha_i} e^{i(\omega_i t + \varphi_i)} + |A_i| e^{\alpha_i} e^{-(j\omega_i + \varphi_i)},$$

d.h. aber

$$h_i(t) = 2|A_i| e^{\alpha_i t} \cos(\omega_i t + \varphi(\omega_i)). \quad (1.116)$$

Die Existenz einer komplexen Nullstelle von $Q(s)$ zusammen mit der mit ihr assoziierten konjugiert komplexen Nullstelle bedeutet also, dass die Impulsantwort $h(t)$ einen Summanden der Form (1.116) enthält. h_i ist eine (Ko-)Sinusschwingung der Frequenz ω_i , die um $\varphi(\omega_i)$ phasenverschoben ist. Zur Zeit $t = 0$ ist die Amplitude durch $2|A_i|$ gegeben. Für

- $\alpha_i < 0$ fällt dann die Amplitude exponentiell gegen Null, und für
- $\alpha_i > 0$ steigt sie exponentiell; das System ist insgesamt instabil, und für
- $\alpha_i = 0$ schwingt das System mit unveränderter Amplitude $|A_i|$ fort.

ω_i ist eine *Eigenfrequenz* des Systems. Offenbar können solche, mit den Eigenfrequenzen verbundenen Eigenschwingungen nur dann auftreten, wenn es mindestens zwei zueinander konjugiert komplexe Nullstellen von $Q(s)$ gibt. Dies bedeutet, dass die homogene Dgl von mindestens 2-ter Ordnung sein muß, soll es überhaupt zu Eigenschwingungen kommen.

Es sei $\alpha_i < 0$. Ist $|\alpha_i|$ klein, so geht $e^{\alpha_i t}$ nur langsam gegen Null und man spricht von *schwacher Dämpfung*. Ist $|\alpha_i|$ so groß, dass es kaum noch zu einer oszillatorischen Bewegung kommt, so hat man es mit einer *starken Dämpfung* zu tun.

Für die Impulsantwort eines Systems n -ter Ordnung gilt nun

$$h(t) = \sum_{i=1}^{n-r} h_i(t).$$

Denn man erhält für jedes Paar zueinander konjugiert komplexer Nullstellen von $Q(s)$ einen Term der Form (1.113). Gibt es r solche Paare, so gibt es r Terme der Form (1.113) und es sind $2r$ Nullstellen absorbiert worden. Die verbleibenden $n - 2r$ Terme müssen von der Form (1.112) sein. Damit gibt es $k = n - 2r + r = n - r$ Terme $h_i(t)$, aus denen sich die Impulsantwort zusammensetzt.

1.2.14 Der Fall mehrfacher Pole

Weist das Polynom $Q(s)$ mehrfache Nullstellen, so ist es übersichtlicher, den Fall ausschließlich reeller mehrfacher Nullstellen und den mehrfachen komplexer Nullstellen getrennt zu betrachten.

Mehrfache reelle Pole Die Nullstelle $s_i \in \mathbb{R}$ des Polynoms $Q(s)$ trete mit der Vielfachheit n_i auf, $n_i \geq 1$, $n_1 + n_2 + \dots + n_k$, $k \leq n$. Dann gilt

$$Q(s) = \prod_{i=1}^n (s - s_i)^{n_i} = \frac{A_1}{(s - s_1)} + \frac{A_2}{(s - s_1)^2} + \dots + \frac{A_{n_1}}{(s - s_1)^{n_1}} + \dots + \frac{L_1}{(s - s_{n_k})} + \dots + \frac{L_{n_k}}{(s - s_k)^{n_k}}. \quad (1.117)$$

Natürlich ist (1.112) ein Spezialfall von (1.117), denn (1.112) geht in (1.117) über, wenn $n_1 = \dots = n_k = 1$. Wieder läßt sich sagen, dass das Gesamtsystem aus der Parallelschaltung von Subsystemen besteht, wobei jedes Subsystem durch eine Transferfunktion der Form $H_{ik}(s) = C/(s - s_i)^k$ mit $k \geq 1$, ist. $C \in \mathbb{R}$ ist eine der Koeffizienten $A_j, \dots, L_j$. H_{ik} läßt sich in der Form $H_{ik}(s) = C[1/(s - s_i)]^k$ schreiben, d.h. sie ist dem k -fachen Produkt der Transferfunktionen $1/(s - s_i)$ proportional. Nach Abschnitt 1.2.10, Gleichung (1.92), p. 40, kann dieser Sachverhalt als Repräsentation eines k -fach rückwirkungsfrei hintereinandergeschalteten Systems mit der Transferfunktion $1/(s - s_i)$ aufgefaßt werden. Da $s_i \in \mathbb{R}$ vorausgesetzt wurde sind diese letzteren Komponenten Systeme erster Ordnung. Das Gesamtsystem kann also als eine Parallelschaltung von Teilsystemen, die aus rückwirkungsfrei hintereinandergeschalteten Systemen erster Ordnung bestehen, beschrieben werden.

Es soll noch die inverse Laplacetransformation der Komponenten betrachtet werden; es ist

$$L^{-1} \left(\frac{A_k}{(s - s_i)^k} \right) = A_{ik} t^k e^{s_i t}. \quad (1.118)$$

Die Komponente $h_i(t)$ der Impulsantwort, die durch s_i definiert ist, hat also die Form

$$h_i(t) = \sum_{k=1}^n \sum_{i=1}^n A_{ik} t^k e^{s_i t} = e^{s_i t} \sum_{i=1}^n A_{ik} t^k; \quad (1.119)$$

Der Ausdruck $t^k e^{s_i t}$ ist als Gamma-Funktion bekannt und ergibt sich als k -fache Faltung der Exponentialfunktion $e^{s_i t}$. Diese Impulsantwort ergibt sich, wenn man k Systeme mit der Impulsantwort $e^{s_i t}$ rückwirkungsfrei hintereinanderschaltet; der Ausgang des einen Systems dient als Eingang für das folgende. (1.117) entspricht *im Zeitbereich* der aus (1.114) hergeleiteten Aussage.

Mehrfache komplexe Nullstellen Es sei nun s_i eine k -fache, $k > 1$, komplexe Nullstelle. Das Polynom $Q(s)$ muß dann das Teilprodukt

$$\sum_{j=1}^n (s - s_i)^j$$

enthalten, dem in der Partialbruchzerlegung von $H(s)$ der Ausdruck

$$\frac{D_1 s_i + E_1}{s_i^2 + p_1 s_i + q_1} + \frac{D_2 s_i + E_2}{(s_i^2 + p_1 s_i + q_1)^2} + \dots + \frac{D_k s_i + E_k}{(s_i^2 + p_1 s_i + q_1)^k} \quad (1.120)$$

entspricht. Da der komplexen Zahl s_i eine konjugiert komplexe Zahl als Nullstelle von $Q(s)$ entspricht, liefert die inverse Laplacetransformation den Ausdruck

$$h_i(t) = A_i t^k e^{\alpha_i t} [\cos(\omega_i t + \varphi_i)], \quad A_i \in \mathbb{R} \quad (1.121)$$

als i -ten Summanden für die Impulsantwort $h(t)$.

Ist also die Transferfunktion oder die Systemantwort bekannt, so liefern die korrespondierenden Summanden, die durch die Partialbruchzerlegung entstehen, die entsprechenden Summanden der Impulsantwort. Die Methoden der Berechnung der reellen Parameter A_i , α_i etc brauchen hier nicht angegeben zu werden, da sie bei einer Analyse des visuellen System bestenfalls als freie Parameter geschätzt werden.

1.2.15 Elementare Systemelemente

In Abschnitt 1.2.10, (vergl. Gleichung (1.95), p. 40) wurde gezeigt, dass eine gegebene Transferfunktion $H(s)$ als Summe von Funktionen $H_i(s)$ dargestellt werden kann, die parallelgeschaltete Systeme repräsentieren; die $H_i(s)$ können wiederum aus rückwirkungsfrei hintereinandergeschalteten Systemen bestehen. Es sollen hier einige elementare Typen von Systemen betrachtet werden, die durch die $H_i(s)$ repräsentiert werden.

1.2.16 Elementare Übertragungsglieder

Das P-Glied Es sei $H(s)$ die Transferfunktion eines Systems. Gilt $H(s) = V \in \mathbb{R}$ eine Konstante, so heißt das System ein *Proportional-Glied*.

Beispiel 1.2.10 In einem elektrischen Schaltkreis sind die Spannung U , der Strom I und der Widerstand R sind gemäß der Gleichung $U = RI$ miteinander verknüpft. Betrachtet man die Spannung U als Eingangsgröße und I als Ausgangsgröße, so ist

$$x(t) = vu(t), \quad v = 1/R \quad (1.122)$$

mit $x(t) = I(t)$, $u(t) = U(t)$. Die Transferfunktion ist durch

$$G(s) = H(s)F(s) = vF(s) \quad (1.123)$$

gegeben; im genannten Beispiel ist $v = 1/R$ eine reelle Konstante. Ein solches Glied muß offenbar nicht durch eine Differentialgleichung beschrieben werden. $\square$

Das Integrier-Glied (I - Glied) Hier soll gelten

$$x(t) = \int_0^t u(\tau)h(t - \tau) d\tau = \int_0^t u(\tau) d\tau \quad (1.124)$$

d.h. das Ausgangssignal soll dem Integral des Eingangssignals sein. Daraus folgt $h(t - \tau) = 1$. Bildet man die Laplacetransformierte von h , so hat man

$$L(h(t)) = \int_0^\infty 1e^{-st} dt = \frac{1}{s},$$

d.h. die Transferfunktion ist in diesem Fall durch

$$H(s) = \frac{1}{s} \quad (1.125)$$

gegeben.

Beispiel 1.2.11 Integrator Betrachtet man die Wechselspannung $U(t)$ an einem Kondensator mit der Kapazität C sowie den Strom $I(t)$, so gilt die Beziehung $dU(t)/dt = I(t)/C$ oder

$$U(t) = \frac{1}{C} \int_0^t I(\tau) d\tau$$

Der Kondensator wirkt wie ein *Integrator*. □

Beispiel 1.2.12 Ein Neuron "feuert" dann ein Aktionspotential (Spike), wenn das Membranpotential $V(t)$ einen bestimmten Schwellenwert V_0 erreicht hat. Das Membranpotential wird durch erregende Impulse oder gradierte Potentiale aufgebaut. Bezeichnet man mit $u(t)$ das erregende Potential, so hat man

$$V(t) = x(t) = \int_0^t u(\tau) d\tau; \quad (1.126)$$

für $V(t_0) = V_0$ wird der Spike getriggert und das Membranpotential V wieder auf Null gesetzt. Diese Gleichung charakterisiert den *perfekten Integrator*.

Ist $u(\tau) = c$ eine Konstante, so ist die *Feuerrate*, d.h. die Rate, mit der Spikes erzeugt werden, durch $r = 1/t_0 = c/V_0$ gegeben, d.h. r ist proportional dem Input $u = c$.

Bei einem anderen Integratortyp werden die Input-Effekte linear summiert und zerfallen exponentiell, so lange das Membranpotential unterschwellig bleibt. Erreicht $V(t)$ den Wert V_0 , so wird der Spike getriggert und $V(t) = 0$ gesetzt. Es gilt

$$x(t) = V(t) = \int_0^t u(\tau) e^{-\lambda(t-\tau)} d\tau \quad (1.127)$$

für $0 < V(t) < V_0$. Ist insbesondere $u(t) = k$ eine Konstante, so erhält man

$$v_1 = k(1 - e^{-at_0})/a,$$

so dass

$$t_0 = -\frac{1}{a} \log \left(1 - v_1 \frac{a}{k} \right).$$

Die Feuerrate $1/t_0 = a/\log(1 - v_1/k)$ ist demnach nicht mehr, wie beim perfekten Integrator, proportional zur Stärke bzw. Intensität k des Eingangssignals. Auch für dieses Modell sind die Reaktionen insbesondere auf stochastische Folgen von Impulsen untersucht worden; eine Zusammenfassung der Resultate findet man in Holden (1977).

Man kann diesen Ausdruck so auffassen, dass das Neuron ein System mit der Gewichtungsfunktion, d.h. mit der Impulsantwort $h(t) = e^{-\lambda t}$ ist.

Zur Erklärung des Ausdrucks "leckender" Integrator (*leaky integrator*) nehme man an, dass der Input $u(\tau)$ ein Impuls $m\delta(\tau)$ sei. Aus (1.126) folgt dann $V(t) = me^{-\lambda t}$, d.h. das Membranpotential fällt exponentiell ab; der Integrator "leckt" gewissermaßen.

Besteht $u(\tau)$ aus einer Folge von Impulsen $m\delta(\tau - t_i)$, wobei die Zeitpunkte zufällig sind (etwa durch einen Poisson-Prozeß gegeben sind), so werden $V(t)$ und damit die Folge der getriggerten Spikes ebenfalls Zufallsprozesse sein. Je größer dabei die Abstände zwischen den erregenden Impulsen sind, desto mehr wird $V(t)$ zwischen den Impulsen absinken und das Erreichen von V_0 wird unwahrscheinlicher; die Spikedichte nimmt ab.

Der Frequenzgang ist durch

$$H(j\omega) = \int_0^t h(t)e^{-j\omega t} dt = \frac{1}{\lambda + j\omega} \quad (1.128)$$

gegeben. Es ist

$$H(j\omega) = \frac{1}{\lambda + j\omega} = \frac{\lambda - j\omega}{(\lambda + j\omega)(\lambda - j\omega)},$$

d.h.

$$H(j\omega) = \frac{\lambda}{\lambda^2 + \omega^2} - i \frac{\omega}{\lambda^2 + \omega^2},$$

woraus

$$|H(j\omega)| = \frac{1}{\lambda^2 + \omega^2}.$$

folgt. Hieraus folgt, dass die Übertragung für $\omega = 0$ maximal ist; die Amplitude $|H(j\omega)|$ fällt monoton mit ω . Der Integrator mit Leck ist also ein Tiefpaßfilter.

In Abschnitt 1.2.5, Beispiel 1.2.4, p. 21, wurde eine Dgl erster Ordnung betrachtet; ihre Lösung für einen Impuls als Eingang ist durch $k \exp(-at)$ gegeben. Dies ist die Impulsantwort. Daraus folgt, dass der Integrator durch eine Dgl erster Ordnung beschrieben werden kann. $\square$

Das Differenzier-Glied (D - Glied): Hier wird gefordert, dass

$$x(t) = \frac{du(t)}{dt} \quad (1.129)$$

gilt. Ist etwa $x(t) = U(t)$ die Spannung an einer Induktivität der Größe L und $I(t)$ der Strom, so gilt $U(t) = LdI(t)/dt$. Die Transferfunktion $H(s)$ ist durch

$$H(s) = Ls \quad (1.130)$$

gegeben.

Viele Systeme können als aus P -, I - und D -Gliedern zusammengesetzt gedacht werden. Bei einer solchen Zusammensetzung von Bauteilen muß im allgemeinen die Differentialgleichung für das gesamte System bestimmt werden, wenn man das Übertragungsverhalten des Systems bestimmen will. Ein interessanter Spezialfall ist aber die rückwirkungsfreie Zusammenschaltung. Diese gelingt, wenn man etwa zwischen die Systeme einen Verstärker mit dem Verstärkungsfaktor 1, einem hohen Eingangswiderstand und einem geringen Ausgangswiderstand schaltet; eine nähere Diskussion findet man in Varju (1977).

1.2.17 Teilsysteme zweiter Ordnung

In Abschnitt 1.2.5 wurde gezeigt, dass einfache komplexe Pole der Transferfunktion implizieren, dass die Impulsantwort Summanden der Form

$$h_i(t) = 2|A_i|e^{\alpha_i t} \cos(\omega_i t + \varphi(\omega_i)) \quad (1.131)$$

enthält. Wird also das System mit einem Impuls erregt, so enthält die Impulsantwort eine mit der Frequenz ω_i schwingende Komponente, deren Amplitude für $\alpha_i < 0$ exponentiell abklingt. Die durch $h_i(t)$ repräsentierte Komponente entspricht einem Teilsystem 2-ter Ordnung, das wiederum als aus zwei rückwirkungsfrei hintereinandergeschalteten Systemen 1-ter Ordnung zusammengesetzt gedacht werden kann. Um dies zu sehen, sei das erste System durch die Dgl 1-ter Ordnung $\tau_1 \dot{y} + \tau_2 y = u(t)$, und das zweite System durch die Dgl $\sigma_1 \dot{g} + \sigma_2 g = y$ charakterisiert. Die Störung des ersten Systems ist die Eingangsfunktion $u(t)$, und die Störung für das zweite System ist die Antwort y des ersten Systems. Beide Systeme können zusammengefaßt werden, indem man die zweite Dgl noch einmal differenziert und das Resultat sowie die zweite Dgl in die erste einsetzt. Es ist $\sigma_1 \ddot{g} + \sigma_2 \dot{g} = \dot{y}$, und in die erste Dgl eingesetzt ergibt sich $\tau_1(\sigma_1 \ddot{g} + \sigma_2 \dot{g}) + \tau_2(\sigma_1 \dot{g} + \sigma_2 g) = u$, d.h.

$$\tau_1 \sigma_1 \ddot{g} + (\tau_1 \sigma_2 + \tau_2 \sigma_1) \dot{g} + \tau_2 \sigma_2 g = u, \quad (1.132)$$

also eine Dgl 2-ter Ordnung für die Funktion $g(t)$. Diese Dgl entspricht offenbar der Dgl (1.49), p. 26, wo sie bei der Betrachtung eines harmonischen Oszillators mit Dämpfung hergeleitet wurde. Dort, d.h. in Beispiel 1.2.5, wurde auch gezeigt unter welchen Bedingungen dieser Dgl eine Impulsantwort der Form (1.130) entspricht. Setzt man, um die Notation einfach zu halten,

$$a_0 = \sigma_2 \tau_2, \quad a_1 = \sigma_1 \tau_2 + \sigma_2 \tau_1, \quad a_2 = \sigma_1 \tau_1$$

so ergibt sich für das charakteristische Polynom der DGL (1.132)

$$a_2 \lambda^2 + a_1 \lambda + a_0 = 0 \quad (1.133)$$

Die Wurzeln des charakteristischen Polynoms für diese Gleichung sind durch

$$\lambda_{1,2} = -\frac{a_1}{2a_2} \pm \frac{1}{2a_2} \sqrt{a_1^2 - 4a_0a_2}.$$

gegeben; ob die Impulsantwort die Form (1.131) annimmt, hängt von der Relation zwischen den Koeffizienten a_0 , a_1 und a_2 ab. Wir betrachten hier noch einmal den Fall, dass tatsächlich (1.131) gilt. Damit diese Beziehung gilt, müssen λ_1 und λ_2 komplex sein, und dies ist der Fall, wenn $a_1^2 - 4a_0a_2 < 0$ ist. Man kann demnach

$$\lambda_{1,2} = -\frac{a_1}{2a_2} \pm \frac{i}{2a_2} \sqrt{4a_0a_2 - a_1^2} \quad (1.134)$$

schreiben (wird der Radikand negativ, so wird die Wurzel imaginär und Werte $\lambda_{1,2}$ werden reell). Demnach wäre in (1.131) $\alpha_i = -a_1/2a_2$ und ω_i ist durch die Wurzel

definiert. Die Schwingung klingt also nur dann exponentiell ab, wenn $0 < a_1$; demnach heißt a_1 die *aktuelle Dämpfungskonstante*. Ist $a_1 = 0$, so ist die Schwingung ungedämpft, und ω_i ist durch $\omega_{0i} = \sqrt{a_0/a_2}$ gegeben; ω_{0i} ist die *Eigenfrequenz des ungedämpften Systems*. Es sei $\tilde{a}_1^2 = 4a_0a_2$; für $a_1 = \tilde{a}_1$ verschwindet die Wurzel und es folgt $\omega_i = 0$; dann tritt keine Schwingung mehr auf. Deshalb heißt $\tilde{a}_1$ die *kritische Dämpfung*. Das Verhältnis von aktueller zu kritischer Dämpfungskonstante ist die *Dämpfungskonstante* ζ , d.h. also $\zeta = a_1/\tilde{a}_1 = a_1/(4a_0a_2)$. Man kann nun das Schwingungsverhalten der Lösung g von (1.132) über ω_{0i} und ζ ausdrücken. Aus (1.133) folgt nach Division durch a_0

$$\frac{a_2}{a_0}\lambda^2 + \frac{a_1}{a_0}\lambda + 1 = 0. \quad (1.135)$$

Aus der Definition von ω_{0i} folgt $a_2/a_1 = 1/\omega_0$, und weiter hat man $a_1/a_0 = 2\zeta/\omega_{0i}$, so dass (1.135) in der Form

$$\frac{1}{\omega_{0i}^2}\lambda^2 + \frac{2\zeta}{\omega_{0i}}\lambda + 1 = 0$$

geschrieben werden kann. Multiplikation mit ω_{0i}^2 liefert dann

$$\lambda^2 + 2\zeta\omega_{0i}\lambda + \omega_{0i}^2 = 0. \quad (1.136)$$

Dies ist die charakteristische Gleichung des Systems 2-ter Ordnung in *Standardform*. Der Ausdruck für $\lambda_{1,2}$ nimmt dann die Form

$$\lambda_{1,2} = \alpha_i \pm j\omega_i = -\zeta\omega_{0i} \pm j\omega_{0i}\sqrt{1-\zeta^2} \quad (1.137)$$

an. Man kann hieraus den Effekt der Dämpfung ζ auf ω_i direkt ablesen: es ist ja $\omega_i = \omega_{0i}\sqrt{1-\zeta^2}$; für $\zeta \rightarrow 0$ folgt $\omega_i \rightarrow \omega_{0i}$, und für $\zeta \rightarrow 1$ folgt $\omega_i \rightarrow 0$.

Da in (1.132) keine Ableitungen des Eingangssignales $u(t)$ eingehen, läßt sich der Frequenzgang $H(j\omega)$ des Systems sofort angeben; es ist, wenn man von (1.136) ausgeht,

$$H(j\omega) = \frac{1}{(j\omega)^2 + 2\zeta\omega_{0i}(j\omega) + \omega_{0i}^2} = \frac{1}{\omega_{0i}^2 - \omega^2 + j2\zeta\omega_{0i}\omega}. \quad (1.138)$$

Daraus folgt

$$A(\omega) = |H(j\omega)| = \frac{\sqrt{(\omega_{0i}^2 - \omega^2)^2 + (2\zeta\omega_{0i}\omega)^2}}{(\omega_{0i}^2 - \omega^2)^2 - (2\zeta\omega_{0i}\omega)^2}. \quad (1.139)$$

Man kann nachrechnen, dass $A(\omega)$ ein Maximum annimmt für

$$\omega_{mi} = \omega_{0i}\sqrt{1-\zeta^2}. \quad (1.140)$$

ω_{mi} ist dann die Dämpfung für $\zeta \neq 0$; für $\zeta = 0$ ist dann $\omega_{mi} = \omega_{0i}$. Das Teilsystem mit der Impulsantwort (1.131) wird also bei einem sinusförmigen Eingangssignal um so stärker reagieren, je näher der Wert von ω bei ω_{mi} liegt; dieses ist das Phänomen der *Resonanz*.

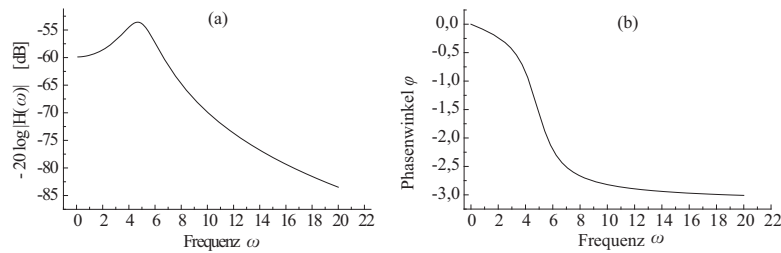
Da Feedback-Systeme als rückwirkungsfreie Hintereinderschaltung von Systemen aufgefaßt werden können, folgt, dass Subsysteme 2-ter Ordnung möglicherweise durch Feedback-Kreise entstehen können und Feedback-Kreise erzeugen also möglicherweise Resonanzphänomene.

1.2.18 Graphische Repräsentationen und Maßeinheiten

Gegeben sei ein System mit dem Frequenzgang $H(j\omega) = H_1(j\omega) + iH_2(j\omega)$. Graphische Darstellungen erlauben gelegentlich ein Abschätzen von interessierenden Eigenschaften des Systems. Es sollen deshalb kurz einige übliche Darstellungen für $H(j\omega)$ zusammen mit einigen gebräuchlichen Maßeinheiten angegeben werden.

Eine erste Möglichkeit, $H(j\omega)$ graphisch darzustellen, ist die *Ortskurve* (*Nyquist-Diagramm*). $H(j\omega)$ ist im allgemeinen eine komplexe Funktion von ω , d.h. $H(j\omega) = H_1(\omega) + iH_2(\omega)$. Man kann nun die Werte von H_1 und H_2 , die ja beide reelle Funktionen von ω sind, für verschiedene Werte von ω berechnen und dann $H(j\omega)$ im Argand-Diagramm darstellen, in der die Ortskurve des durch $H(j\omega) = 1/(\alpha + j\omega)$ angegeben wird; die x -Achse repräsentiert die Werte von H_1 , die y -Achse die von iH_2 , und $H_1(\omega) = \alpha/(\alpha^2 + \omega^2)$, $H_2(\omega) = \omega/(\alpha^2 + \omega^2)$. In der folgenden, als

Abbildungung 1.7: (a) Amplitudengang für System 2-ter Ordnung, Dämpfung $\zeta = .35$, (b) entsprechender Phasengang



Bode-Diagramm bekannten Darstellung wird der Amplituden- und der Frequenzgang getrennt dargestellt. Es ist $H(j\omega) = A(\omega)e^{i\varphi(\omega)}$, wobei $A(\omega) = |H(j\omega)|$, $\varphi(\omega) = \tan^{-1}[H_2(j\omega)/H_1(j\omega)]$. $A(\omega)$ repräsentiert den Amplitudengang, d.h. die Verstärkung oder Abschwächung (Attenuation) der Amplitude eines sinusförmigen Eingangssignals, und $\varphi(\omega)$ den zugehörigen Phasengang, d.h. die zugehörige Phasenverschiebung. Wir betrachten nun den Logarithmus von $H(j\omega)$. Es ist

$$\log H(j\omega) = \log \left(|H(j\omega)|e^{i\varphi(\omega)} \right) = \log |H(j\omega)| + i\varphi(\omega). \quad (1.141)$$

Das *Bode-Diagramm* (Bode-Plot) von $H(j\omega)$ liefert das Bild von (i) $\log_{10} A(\omega)$ und (ii) des Phasenwinkels $\phi(\omega)$ als Funktion von $\log_{10} \omega$; mit $\log_{10}$ ist, wie üblich, der Logarithmus zur Basis 10 gemeint. Abbildung 1.7 zeigt insbesondere den Amplitudengang eines Systems 2-ter Ordnung, dessen System-Funktion in Gleichung (1.138) angegeben wurde. Dabei ist $f = 2\pi/\omega$, $f_{max} = 5$ [Hz]; der Dämpfungsfaktor ist $\zeta = .35$. In Abbildung 1.7 wird der Phasengang Verhalten des gleichen Systems gezeigt. Beide Abbildungen zusammen bilden das Bode-Diagramm.

Gelegentlich wird die x - bzw. die f -Achse des Diagramms nicht direkt in [Hz], sondern in anderen Einheiten dargestellt. Eine solche Einheit ist die *Dekade*. Die Dekade ist die Länge des Intervalls $(\omega, 10\omega)$, denn es ist $\log_{10}(10\omega) - \log_{10} \omega =$

$\log_{10}(10\omega/\omega) = \log_{10} 10 = 1$. Gelegentlich wird auch von der *Oktave* als Einheit Gebrauch gemacht. Die Oktave ist die Länge des Intervalls $(\omega, 2\omega)$; in der Musik sind zwei Töne mit der Frequenz ω_1 und ω_2 bekanntlich dann eine Oktave voneinander entfernt, wenn $\omega_2/\omega_1 = 2$. Auf der $\log_{10}$ -Achse entspricht der Oktave dementsprechend die Länge $\log_{10}(2\omega) - \log_{10} \omega = \log_{10}(2\omega/\omega) = \log_{10} 2 \approx .3$, so dass 1 Oktave $\approx .3$ Dekaden. Einer Dekade oder einer Oktave entspricht nicht eine konstante *Differenz*, sondern ein konstantes *Verhältnis*.

Die Einheit der *Ordinate* ist im Bode-Diagramm das *Dezibel*; das Dezibel ist der zehnte Teil eines *Bels*, dass durch

$$B = \log_{10} \frac{I}{I_0} \quad (1.142)$$

definiert ist, wobei hier I_0 und I allgemein Intensitäten repräsentieren. Die Definition des Bel entspricht der Fechnerschen Beziehung für die Empfindungsstärke:

$$f(I) = \alpha \log \frac{I}{I_0},$$

wobei α ein Skalenfaktor und I_0 die Absolutschwelle ist. In (1.142) wurde insbesondere der Logarithmus zur Basis 10 und $\alpha = 1$ gewählt. Für $B = k$ erhält man insbesondere

$$\frac{I}{I_0} = 10^k,$$

d.h. $I = 10^k I_0$, die Intensität I unterscheidet sich von I_0 um einen Faktor, der durch k Zehnerpotenzen definiert ist.

Die Bel-Einheit ist häufig zu groß im Verhältnis zu den betrachteten I -Werten; eine feinere Abstufung erhält man, wenn man die verschiedenen I -Werte in bezug auf Zehntel 10-Potenzen diskutiert. Da 10 Zehntel gerade eine Zehnerpotenz ergeben, hat man dementsprechend

$$\text{Dezibel} \stackrel{\text{def}}{=} 10 \log_{10} \frac{I}{I_0} = k \quad [\text{dB}] \quad (1.143)$$

I_0 definiert den (willkürlichen) Nullpunkt der Skala, denn $L = 0$ für $I = I_0$. Es seien I_j und I_k zwei Intensitäten derart, dass

$$10 \log_{10} \frac{I_j}{I_0} = j \text{ [dB]}, \quad 10 \log_{10} \frac{I_k}{I_0} = k \text{ [dB]}, \quad j < k$$

Dann folgt

$$10(\log_{10} I_k - \log_{10} I_j) = 10 \log_{10} \frac{I_k}{I_j} = k - j,$$

und

$$\frac{I_k}{I_j} = 10^{(k-j)/10}.$$

Der Unterschied zwischen I_k und I_j , d.h. der Quotient zwischen I_k und I_j entspricht also gerade $(k - j)/10$ Zehnerpotenzen.

Der Energie einer Frequenz entspricht das Quadrat der Amplitude, also $A^2(\omega)$. Setzt man die Intensität $I(\omega)$ gleich der Energie, also $I(\omega) = A^2(\omega)$, so liefert (1.143)

$$10 \log_{10} \frac{A^2}{A_0^2} = 20 \log_{10} \frac{A}{A_0}, \quad (1.144)$$

wobei der Quotient auf der rechten Seite durch Amplituden definiert ist; dafür ist der Skalenfaktor jetzt 20 statt 10. Die Dezibel-Skala wird oft direkt in der Form des Quotienten auf der rechten Seite von (1.144) eingeführt (z.B. Röhler (1973), p. 27)⁸

Auf der Ordinate wird im Bode-Diagramm die Größe $20 \log_{10} A(\omega)$ bzw. $20 \log_{10} A(f)$, $f = \omega/2\pi$, abgetragen. ω_1 und ω_2 mögen sich gerade um eine Oktave unterscheiden, und es sei $A_1 = A(\omega_1)$, $A_2 = A(\omega_2)$.

$$20 \log_{10} A_2 - 20 \log_{10} A_1 = k.$$

A_1^2 und A_2^2 unterscheiden sich dann um k Zehntel Zehnerpotenzen, die Amplituden wegen

$$\frac{A_2}{A_1} = 10^{k/20}$$

um k Zwanzigstel 10-er Potenzen. Man betrachte nun das Ausgangssignal eines Systems, an dem das Eingangssignal $\sin(\omega x)$ bzw. $\sin(\omega t)$ liegt; das Ausgangssignal ist dann durch $|H(\omega)|(\omega) \sin(\omega x)$ (und entsprechend für t) gegeben. Die Amplitude des Ausgangssignals ist $|H(\omega)|$. Weiter sei $|H(\omega_1)| = 1$. Dann definiert

$$M(\omega) \stackrel{\text{def}}{=} 20 \log_{10} \frac{|H(\omega)|}{|H(\omega_1)|} = 20 \log_{10} |H(\omega)| \quad [\text{dB}] \quad (1.145)$$

die *Dämpfung* (*Gain, Attenuation*) des Systems. Trägt man auf der x -Achse die Frequenzen in Oktaven oder Dekaden ab, so zeigt der Bode-Plot die Dämpfung in pro Oktave oder Dekade.

Um die Beziehung zwischen $A(\omega)$ - und Dezibel-Werten zu veranschaulichen, sollen einige numerische Beispiele gegeben werden. $A(\omega) = .01$ entspricht $20 \log_{10} A(\omega) = -40$ dB. $A(\omega) = .1$ entspricht - 20 dB; Für $A(\omega) = .5$ erhält man - 6 dB; $A(\omega) = 1$ ist gerade 0 dB. Zusammenfassend läßt sich sagen, dass

- (i) der Verdoppelung von $A(\omega)$ eine Erhöhung um ≈ 6 dB entspricht, denn liefert $A(\omega)$ gerade $y = 20 \log_{10} A(\omega)$ dB, so ist $y_1 = 20 \log_{10}(2A(\omega)) = 20 \log_{10} A(\omega) + 20 \log_{10} 2 = 20 \log_{10} A(\omega) + 6.0206 \dots \approx y + 6$ dB.
- (ii) Erhöht man $A(\omega)$ um den Faktor 10, d.h. geht man zu $10 A(\omega)$ über, so bedeutet dies eine Erhöhung um 20 dB. Denn $20 \log_{10}(10 A(\omega)) = 20 \log_{10} A(\omega) + 20 \log_{10} 10$, und $20 \log_{10} 10 = 20$. Einem *Faktor* α entspricht ein Zuwachs von $20 \log_{10} \alpha$.

Im Bode-Diagramm für die Phase wird $\phi(\omega)$ direkt, d.h. nicht logarithmiert, abgetragen; die Maßeinheit ist das *Grad*.

⁸Röhler, R. Biologische Kybernetik. Teubner-Studienbücher Biologie, Stuttgart 1973.

1.2.19 Filter und ihre Charakterisierung

Gegeben sei ein System mit dem Frequenzgang $H(j\omega)$. Dann gibt $A(\omega) = |H(j\omega)|$ die Amplitude der Antwort auf $\sin(\omega t)$ an. Die Frequenz $\omega = 2\pi f$ kann Werte im Bereich $0 \leq \omega < \infty$ annehmen. Es ist aber möglich, dass $A(\omega)$ nur Werte ungleich Null in einem bestimmten Bereich der positiven Achse annimmt; dies führt zu der folgenden

Definition 1.2.10 *Ein System habe den Frequenzgang $H(j\omega)$, und es sei $A(\omega) = |H(j\omega)|$. Es sei $A_0 = \max A(\omega)$, und es $A(\omega) \neq 0$ nur für $\omega_a < \omega < \omega_b$, und $\omega_c \in (\omega_a, \omega_b)$ sei derart, dass $A(\omega_c) = \frac{1}{2}A_0$. Dann heißt das System ein Filter für das Intervall (ω_a, ω_b) mit der Eckfrequenz ω_c . Gilt insbesondere $\omega_a = 0$, so heißt das System Tiefpaßfilter, für $\omega_a > 0$, $\omega_b = \infty$ heißt es Hochpaßfilter und für $0 < \omega_a < \omega_b < \infty$ heißt es Bandpaßfilter. Es sei $\omega_0 = \frac{1}{2}(\omega_b - \omega_a)$; dann heißt ω_0 die Zentralfrequenz (center frequency) des Filters; gilt insbesondere $(\omega_b - \omega_a) \ll \omega_0$, so heißt das System Schmalbandfilter (narrow-band filter)*

Der Begriff des Filters bezieht sich also auf den Frequenzbereich, für den das System durchlässig ist. Die Forderung, dass $A(\omega) = 0$ außerhalb des Intervalles (ω_a, ω_b) ist, wird bei tatsächlich gegebenen Filtern kaum in strenger Form erfüllt sein; sie trifft nur für den *idealen Filter* zu. Das Konzept des Filters soll nun an einigen Beispielen illustriert werden.

Beispiel 1.2.13 (Einpöler Tiefpaß) Die Transferfunktion habe nur einen Pol, d.h. das Polynom $Q(s)$ habe nur eine Nullstelle $s = \alpha$. α ist notwendig reell, denn die Existenz einer komplexen Nullstelle impliziert die Existenz einer zweiten, dazu konjugiert komplexen Nullstelle, im Widerspruch zur Forderung, dass $Q(s)$ eben nur eine Nullstelle hat. Dann läßt sich $H(s)$ in der Form $H(s) = A/(s - \alpha)$ darstellen, $A \in \mathbb{R}$, und die Impulsantwort ist durch

$$h(t) = a e^{\alpha t} \quad (1.146)$$

gegeben. Das System ist stabil, wenn $\alpha < 0$. Es ist $A(\omega) = a/(\omega^2 + \alpha^2)$, $\varphi(\omega) = \tan^{-1}(\omega/\alpha)$. Offenbar geht $A(\omega)$ gegen Null für $\omega \rightarrow \infty$. Da $\omega^2 > 0$ für alle $\omega > 0$ ist $A(\omega)$ maximal für $\omega = 0$; je niedriger der Wert von ω , desto größer die Amplitude der Ausgangsschwingung, und in diesem Sinne ist der Filter ein Tiefpaß.

Betrachten wir nun einen Filter mit dem Frequenzgang $H(j\omega) = (1/(\beta + j\omega))^n$; ihm entspricht die Impulsantwort $h(t) = t^n e^{\beta t}$ (vergl. die Abbildungen A.9, 174, etc. im Anhang). Um $A(\omega) = |H(j\omega)|$ zu bestimmen muß man nur von der Tatsache Gebrauch machen, dass ja $H(j\omega) = |H(j\omega)|e^{i\varphi(\omega)}$ gilt. Es sei nun $H_0(j\omega) = 1/(\beta + j\omega) = |H_0(j\omega)|e^{i\varphi_0(\omega)}$; man hat dann sofort

$$H(j\omega) = |H_0(j\omega)|^n e^{in\varphi_0(\omega)} = \left(\frac{1}{\beta^2 + \omega^2} \right)^{n/2} e^{in\varphi_0(\omega)}. \quad (1.147)$$

Es werde die Veränderung von $A(\omega) = |H(j\omega)| = |H_0(j\omega)|^n$ mit ω für $n > 1$ betrachtet. Für vorgegebenes n fällt $|H(j\omega)|$ monoton mit ω . Von einem gewissen

$\omega = \omega_1$ an wird $1/(\beta^2 + \omega^2) < 1$, und es folgt, dass $|H_0(j\omega)|^n$ um so schneller gegen Null gehen wird, je größer der Wert von n ist. Dies bedeutet, dass der Filter um so undurchlässiger für hohe Frequenzen wird, je mehr Glieder mit dem Frequenzgang $1/(\beta + j\omega)$ rückwirkungsfrei hintereinander geschaltet sind.

Wie im Bode-Diagramm werde nun der Verlauf von $\log_{10} A(\omega)$ betrachtet. In dB-Einheiten hat man $20 \log_{10} A(\omega) = n 20 \log_{10} \sqrt{1/(\beta^2 + \omega^2)} = -n 20 \log_{10} \sqrt{\beta^2 + \omega^2}$. Nun ist von einem gewissen Wert von ω an $\log_{10} \sqrt{\beta^2 + \omega^2} \approx \log_{10} \omega$, d.h. aber

$$M(\omega) = 20 \log_{10} A(\omega) \approx -20n \log_{10} \omega. \quad (1.148)$$

Da bei einem Bode-Diagramm $20 \log_{10} A(\omega)$ gegen $\log_{10} \omega$ aufgetragen wird geht $M(\omega)$ in eine Gerade mit der Steigung $-20n$ dB/Dekade ≈ -6 dB/Oktave über, denn es ist ja 1 Oktave $\approx .3$ Dekaden, und $20 \times .3 \approx 6$. Aus der Steigung dieser Asymptote kann also zumindest im Prinzip der Wert von n geschätzt werden. Es sei angemerkt, dass diese Schätzung von n , also der Anzahl hintereinandergeschalteter Filter, nicht davon abhängt, dass die Filter alle den gleichen Frequenzgang $1/(\beta + j\omega)$ haben. So sei etwa

$$H(j\omega) = \prod_{j=1}^n \frac{1}{\beta_j + j\omega}.$$

Dann kann man stets eine komplexe Zahl H_0 finden derart, dass

$$H(j\omega) = \prod_{j=1}^n \frac{1}{\beta_j + j\omega} = H_0^n, \quad \text{oder} \quad \left(\prod_{j=1}^n \frac{1}{\beta_j + j\omega} \right)^{1/n} = H_0,$$

und es existiert eine Zahl β derart, dass $H_0(j\omega) = 1/(\beta + j\omega)$. Hierauf läßt sich dann wieder (1.147) anwenden. $\square$

Beispiel 1.2.14 (RC-Glieder) RC-Glieder spielen bei der Beschreibung der Aktivierung von Neuronen eine zentrale Rolle: das Membranpotential hängt vom Stromfluß (Ionenfluß) durch die Membran und damit von deren Widerstand sowie von ihrer Kondensatorwirkung ab (Hodgkin und Huxley (1952)). Es gelten die folgenden Grundbeziehungen zwischen der Spannung $v(t)$, dem Strom $I(t)$, der vom Strom $I(t)$ unabhängigen Eingangsspannung $e(t)$ und dem vom dem von der Spannung $v(t)$ unabhängigen Eingangsstrom $i_g(t)$, wobei R die Größe des (Ohmschen) Widerstandes und C die Kapazität des Kondensators ist:

$$v(t) = RI(t) \quad (1.149)$$

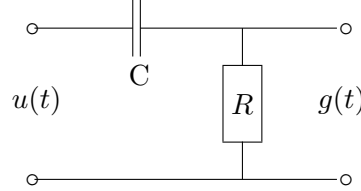
$$v(t) = \frac{1}{C} \int_0^t i(\tau) d\tau + v(0) \quad (1.150)$$

$$I(t) = C \frac{dv(t)}{dt} \quad (1.151)$$

Es werden die folgenden Schaltungen betrachtet:

- a. **Hochpaßfilter** Hier werden der Kondensator und der Widerstand in Reihe geschaltet (Abb. 1.8): $v_r(t)$ ist die Spannung am Widerstand, v_c ist die Span-

Abbildung 1.8: Hochpaßfilter



nung am Kondensator. Nach den Kirchhoffschen Regeln ist die Gesamtspannung $v(t)$ durch die Summe $v(t) = v_c(t) + v_r(t)$ gegeben. $I(t)$ sei der Strom; natürlich muß $I(t) = 0$ sein, wenn $e(t)$ einen konstanten Wert hat. Jedenfalls muß nach (1.149)

$$v_r(t) = RI(t),$$

und

$$v_c(t) = \frac{1}{C} \int_0^t I(\tau) d\tau + v_c(0)$$

gelten, wobei $v_c(0)$ die Spannung am Kondensator zur Zeit $t = 0$ ist. Also folgt einerseits

$$e(t) = v_c(t) + v_r(t) = v_c(t) + RI(t) = v_c(t) + RC \frac{dv_c(t)}{dt},$$

nach (1.151), oder

$$\frac{dv_c(t)}{dt} = -\frac{1}{RC}v_c(t) + \frac{1}{RC}e(t) = -a v_c(t) + a e(t) \quad (1.152)$$

wobei $a = 1/RC$ die Zeitkonstante des Systems ist. Hier entspricht $v_c(t)$ der Funktion $x(t)$ und $ae(t)$ der Funktion $u(t)$ in Beispiel 1.2.4, p. 21. Nach (1.22), p. 1.22 ergibt sich für die partikuläre Lösung, mit $v(0) = 0$,

$$v_c(t) = a \int_0^t e^{-a(t-\tau)} e(\tau) d\tau. \quad (1.153)$$

Ist insbesondere $e(t)$ ein Impuls der Stärke c , so erhält man die Impulsantwort $h_c(t) = m a e^{-at}$. Im Zeitpunkt $t = 0$ springt also die Spannung auf $h_c(0) = m a$.

Die Fouriertransformierte von $h_c(t)$ liefert die Systemfunktion bezüglich der Spannung v_c . Es ist

$$F(h_c) = H_c(j\omega) = m a \int_0^\infty e^{-(a+j\omega)t} dt = \frac{m a}{a + j\omega} \quad (1.154)$$

und

$$|H_c(j\omega)| = \frac{m a}{\sqrt{a^2 + \omega^2}}. \quad (1.155)$$

Für wachsende Frequenz $\omega = 2\pi f$ geht demnach die Spannung $v_c(t)$ am Kondensator gegen Null.

Natürlich kann das Verhalten des Systems auch in bezug auf den Strom $I(t)$ beschrieben werden. Aus $e(t) = v_c(t) + v_r(t)$ folgt

$$e(t) = \frac{1}{C} \int_0^t i(\tau) d\tau + RI(t) + v_c(0),$$

und nach Differentiation erhält man

$$\begin{aligned} \frac{de(t)}{dt} &= \frac{1}{C}I(t) + R\frac{di(t)}{dt}, \\ \frac{di(t)}{dt} &= -\frac{1}{RC}I(t) + \frac{1}{R}\frac{de(t)}{dt}. \end{aligned} \quad (1.156)$$

und nach Division durch R

$$\frac{di(t)}{dt} = -\frac{1}{RC}I(t) + \frac{1}{R}\frac{de(t)}{dt} \quad (1.157)$$

Hier ist $x(t) = I(t)$ und $u(t) = \dot{e}(t)/R$, so dass man für die partikuläre Lösung

$$I(t) = \frac{1}{R} \int_0^t e^{-a(t-\tau)} \dot{e}(\tau) d\tau \quad (1.158)$$

erhält. Es sei $i(0) = 0$ und $e(t) = m\delta(t)$ sei ein Spannungsimpuls. Dann muß man erklären, was unter $\dot{e}(t) = m\dot{\delta}(t)$, d.h. unter der Ableitung von $\delta(t)$ verstanden werden soll. Da sich $\delta(t)$ durch $\int f(\tau)\delta(t-\tau)d\tau = f(t)$ definieren läßt, kann man analog für $\dot{\delta}$ vorgehen (Papoulis (1981), p. 13):

$$d\left(\int f(\tau)\delta(t-\tau)d\tau\right)/dt = \int f'(t-\tau)\delta(\tau)d\tau = \int f(\tau)\dot{\delta}(t-\tau)d\tau = f'(t),$$

d.h. $\dot{\delta}(t)$ liefert $f'(t) = df(t)/dt$. Da $d(e^{-at})/dt = -ae^{-at}$, erhält man mithin

$$I_{imp}(t) = \frac{a}{R}e^{-at}. \quad (1.159)$$

Der Spannungsimpuls erzeugt einen exponentiell abklingenden Strom. Die Systemfunktion ergibt sich aus der Fouriertransformierten von $I_{imp}(t)$: es ist

$$F(I_{imp}(t)) = H_i(j\omega) = \frac{a}{R} \int_0^\infty e^{-(a+j\omega)t} dt = \frac{1}{R} \frac{a}{a+j\omega}, \quad (1.160)$$

so dass

$$|H_i(j\omega)| = \frac{1}{R} \frac{a}{\sqrt{a^2 + \omega^2}}, \quad (1.161)$$

und man sieht, dass bei sinusförmigem Eingangssignal die Amplitude des Stroms $I(t)$ ebenfalls gegen Null strebt.

Es sei insbesondere $e(t) = k$ eine konstante Spannung. Dann ist aber $\dot{e}(t) = de(t)/dt = 0$ und aus (1.158) folgt sofort $I(t) \equiv 0$; bei einer Gleichspannung kann ja durch einen Kondensator kein Strom fließen.

Für die Spannung v_r am Widerstand folgt

$$v_r(t) = e(t) - v_c(t). \quad (1.162)$$

Ist $e(t)$ ein Spannungsimpuls, d.h. ist $e(t) = \delta(t)$, so findet man die Impulsantwort $h_r(t)$ für $v_r(t)$ durch Einsetzen:

$$h_r(t) = \delta(t) - h_c(t) = \delta(t) - ae^{-at}. \quad (1.163)$$

Die Fouriertransformierte von h_r ist dann

$$H_r(j\omega) = \int_{-\infty}^{\infty} h_r(t)e^{-j\omega t} dt = \int_{-\infty}^{\infty} \delta(t)e^{-j\omega t} dt - a \int_{-\infty}^{\infty} e^{-(a+j\omega)t} dt$$

d.h.

$$H_r(j\omega) = 1 - \frac{a}{a + j\omega} = -\frac{j\omega}{a + j\omega} \quad (1.164)$$

so dass

$$|H_r(j\omega)| = \frac{\omega}{\sqrt{a^2 + \omega^2}} = \frac{1}{\sqrt{1 + a^2/\omega^2}} \quad (1.165)$$

Für $\omega \rightarrow \infty$ strebt $|H_r(j\omega)|$ offenbar gegen 1, für $\omega \rightarrow 0$ aber gegen Null. Ist also $e(t) = \sin(\omega t)$ eine sinusförmig modulierte Spannung, so ist die Spannung am Widerstand zwar auch sinusförmig (wie bei allen linearen Systemen), allerdings ist die Amplitude dieser Spannung für kleine Frequenzen gering, um dann mit wachsendem ω einem oberen Grenzwert zuzustreben. In bezug auf $v_r(t)$ ist das RC-Glied ein *Hochpaßfilter*. Für steigende Frequenz ω geht die Spannung am Kondensator gegen Null, ebenso der Strom $I(t)$ im Stromkreis, die Spannung am Widerstand strebt aber gegen einen endlichen, positiven Wert.

- b. **Tiefpaßfilter** Hier sind Kondensator und Widerstand parallel geschaltet. für den Strom gilt nach Kirchhoff $I(t) = I_c(t) + I_r(t)$. Ist $v_c(t)$ die Spannung am Kondensator, so folgt $I(t) = Cdv_c(t)/dt$. Offenbar ist hier $e(t) = v_r(t) = v_c(t) = v(t)$, und $v(t) = RI(t)$. Also muß

$$\frac{1}{C}I(t) = \frac{dv(t)}{dt} + \frac{1}{RC}v(t) \quad (1.166)$$

$$\frac{dv(t)}{dt} = -\frac{1}{RC}v(t) + \frac{1}{C}I(t) \quad (1.167)$$

Mit $a = 1/RC$ und $u(t) = I(t)/C$ ist also wieder

$$v(t) = \frac{1}{C} \int_0^t e^{-a(t-\tau)} I(\tau) d\tau$$

Für einen Stromimpuls $I(t) = \delta(t)$ erhält man dann für die Spannung $e(t) = v(t)$ am Widerstand die Impulsantwort

$$h_r(t) = \frac{1}{C}e^{-at}. \quad (1.168)$$

Die Fouriertransformierte von h_r ist

$$H_r(j\omega) = \frac{1}{C} \int_0^\infty e^{-(a+j\omega)t} dt = \frac{1}{C} \left(\frac{1}{a+j\omega} \right), \quad (1.169)$$

so dass für $A(\omega) = |H_r(j\omega)|$ folgt

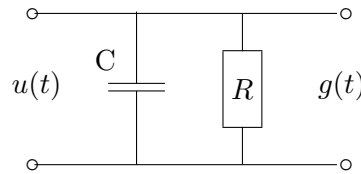
$$|H_r(j\omega)| = \sqrt{\frac{1}{C} \left(\frac{1}{a^2 + \omega^2} \right)}. \quad (1.170)$$

Für $\omega = 0$ ist

$$A(\omega) = |H_r(j\omega)| = \frac{1}{C} \frac{1}{a} = R \neq 0,$$

und für $\omega \rightarrow \infty$ strebt $A(\omega)$ gegen Null. In bezug auf $e(t) = v(t)$ ist der Filter also ein *Tiefpaßfilter*, vergl. Abb. 1.9.

Abbildung 1.9: Tiefpaßfilter



- c. **Hintereinanderschaltung von Tief- und Hochpaßfilter** Schaltet man einen Tief- und einen Hochpaßfilter rückwirkungsfrei hintereinander, so ist die Systemfunktion durch

$$H(j\omega) = -\frac{1}{C} \left(\frac{j\omega}{a+j\omega} \right) \left(\frac{1}{a+j\omega} \right) = -\frac{1}{C} \left(\frac{j\omega}{a^2 - \omega^2 + i2a\omega} \right) \quad (1.171)$$

gegeben. Dann ist

$$|H(j\omega)| = \frac{1}{\sqrt{(a^2 - \omega^2)^2 - 4a^2\omega^2}} \quad (1.172)$$

und

$$\phi(\omega) = \tan^{-1} \left(\frac{\omega(a^2 - \omega^2)}{2a\omega^2} \right). \quad (1.173)$$

□

1.2.20 Maximale Antwort linearer Filter; Matched Filter

Nach Satz 1.2.9, Seite 36, kann die Antwort $g(t)$ zur Zeit t eines Filters durch

$$g(t) = \int_{-\infty}^{\infty} H(j\omega) S(j\omega) e^{j\omega t} d\omega \quad (1.174)$$

dargestellt werden, wobei H die Systemfunktion des Filters und S die Fouriertransformierte des Eingangssignals ist; H ist die Fouriertransformierte der Impulsantwort h des Filters. Es werden zwei Fragen diskutiert:

1. Gegeben sei ein Filter mit der Systemfunktion G und ein Eingangssignal mit der Fouriertransformierten S . Läßt sich eine obere Grenze für die Antwort g angeben?
2. Läßt sich ein Eingangssignal mit der Fouriertransformierten S angeben, für das der durch G charakterisierte Filter eine maximalmögliche Antwort, etwa zu einer Zeit t_0 gibt?

Beide Fragen lassen sich über die Schwartzsche Ungleichung beantworten, derzufolge für zwei beliebige (soll heißen: reelle oder komplexe) Funktionen f_1 und f_2

$$\left| \int_{-\infty}^{\infty} f_1(x) f_2^*(x) dx \right|^2 \leq \int_{-\infty}^{\infty} |f_1(x)|^2 dx \int_{-\infty}^{\infty} |f_2(x)|^2 dx \quad (1.175)$$

gilt. Das Gleichheitszeichen gilt dann, wenn $f_1 = k f_2^*$, $k \in \mathbb{R}$. Auf (1.174) angewendet findet man sofort

$$\left| \int_{-\infty}^{\infty} H(j\omega) S(j\omega) e^{j\omega t} d\omega \right|^2 \leq \int_{-\infty}^{\infty} |H(j\omega)|^2 d\omega \int_{-\infty}^{\infty} |S(j\omega) e^{j\omega t}|^2 d\omega. \quad (1.176)$$

Damit hat man mit der rechten Seite eine obere Schranke für $|g(t)|^2$ und damit auch für $g(t)$, für alle $t \geq 0$ unter der Voraussetzung, dass die Integrale auf der rechten Seite existieren.

Für den Spezialfall $f_1 = k f_2^*$ folgt dann

$$k^2 \left| \int_{-\infty}^{\infty} f_1(x) f_1^*(x) dx \right|^2 \leq k^2 \int_{-\infty}^{\infty} |f_1(x)|^2 dx \int_{-\infty}^{\infty} |f_1(x)|^2 dx$$

und k^2 kürzt sich heraus. Aber $f_1 f_1^* = |f_1|^2$, so dass

$$\left| \int_{-\infty}^{\infty} |f_1(x)|^2 dx \right|^2 \leq \left| \int_{-\infty}^{\infty} |f_1(x)|^2 dx \right|^2;$$

und natürlich gilt nun das Gleichheitszeichen.

Man kann nun den Fall $f_1 = k f_2^*$ auf (1.176) anwenden, d.h. man betrachtet den Fall

$$H(j\omega) = k S^*(j\omega) e^{-j\omega t_0}. \quad (1.177)$$

Dann hat man

$$|g(t_0)|^2 = \left| \int_{-\infty}^{\infty} S(j\omega) S^*(j\omega) d\omega \right|^2 = k \left| \int_{-\infty}^{\infty} |S(j\omega)|^2 d\omega \right|^2 \quad (1.178)$$

d.h. man hat

$$g(t_0) = \max_t g(t) = kE, \quad (1.179)$$

wobei

$$E = \frac{1}{2\pi} \int_{-\infty}^{\infty} |S(j\omega)|^2 d\omega \quad (1.180)$$

die Energie des Signals ist. Die maximale Antwort (zum Zeitpunkt t_0) ist proportional zur Energie des Signals.

Die interessante Bedingung dafür, dass g zu einem Zeitpunkt t_0 die maximal mögliche Antwort ist, ist die Bedingung (1.177). Diese Bedingung ist gleichbedeutend mit

$$S(j\omega) = \frac{1}{k} H^*(j\omega) e^{j\omega t_0}; \quad (1.181)$$

für einen gegebenen Filter mit der Systemfunktion H wird die Antwort maximal (in t_0), wenn das Signal durch eine Fouriertransformierte S definiert ist, die der Bedingung (1.181) entspricht. Man kann diesen Sachverhalt auch umgekehrt ausdrücken:

Definition 1.2.11 Gegeben sei ein Signal – ein Stimulus – mit der Fouriertransformierten $S(j\omega)$. Der Filter, der maximal auf das Signal reagiert, ist durch (1.177) gegeben: $H(j\omega)$ muß proportional zu $S^*(j\omega) e^{-j\omega t_0}$ gewählt werden. Dieser Filter heißt auch Angepasster Filter oder Matched Filter für das Signal.

Es sei s das Signal (der Stimulus), und S die Fouriertransformierte von s . Nach Parsevals Satz (A.74), Seite 170, folgt aber mit $s = f_1$ und $s = f_1$

$$\int_{-\infty}^{\infty} s(t) s(t)^* dt = \int_{-\infty}^{\infty} |s(t)|^2 dt = \int_{-\infty}^{\infty} S S^* d\omega = \int_{-\infty}^{\infty} |S(j\omega)|^2 d\omega \quad (1.182)$$

Nun ist generell die Antwort g des Filters durch die Faltung

$$g(t) = \int_0^t s(\tau) h(t - \tau) d\tau$$

gegeben. Die Fouriertransformierte von h sei H und erfülle die Bedingung (1.181), d.h.

$$H(j\omega) = k S^*(j\omega) e^{-j\omega t}.$$

Dies ist aber die Fouriertransformierte von $s(t - \tau)$: mit $u = t - \tau$ und dementsprechend $\tau = t - u$ hat man

$$\int_0^{\infty} s(t - \tau) e^{-j\omega \tau} d\tau = - \int_{-\infty}^0 s(u) e^{-j\omega(t-u)} du = e^{-j\omega t} \int_0^{\infty} s(u) e^{j\omega u} du = e^{-j\omega t} S^*(j\omega),$$

so dass die Impulsfunktion h des an s angepassten Filters durch

$$h(t) = s(t - \tau) \quad (1.183)$$

gegeben ist.

Die Definition des Matched Filters überträgt sich auf 2-dimensionale Signale, so dass man angepasste Filter für 2-dimensionale Stimulismuster definieren kann.

1.2.21 Energie und Leistung

In vielen Anwendungen spielen die Begriffe der Energie und der Leistung (*power*) bei der Charakterisierung oder der Entdeckung von Signalen eine Rolle. Die Begriffe sollen kurz eingeführt werden.

Der Begriff der Energie wird zunächst im Rahmen der Mechanik eingeführt und ist an den der Kraft gekoppelt. Ist $x(t)$ die Position eines Teilchens zur Zeit t , so ist $v(t) = dx(t)/dt$ seine Geschwindigkeit zur Zeit t . Die Veränderung der Geschwindigkeit ist die Beschleunigung $b(t) = d^2x(t)/dt^2$. Nach Newton ist die auf einen Körper einwirkende Kraft F proportional zu seiner Beschleunigung,

$$F = m b, \quad (1.184)$$

wobei die Proportionalitätskonstante m die Masse des Körpers repräsentiert. Bewegt man einen Körper von x_1 nach x_2 , so verrichtet man Arbeit W , die als Kraft $\times$ Weg definiert ist:

$$W = F x = (mb) x, \quad x = |x_1 - x_2|. \quad (1.185)$$

Für den Weg $\Delta x = x/n$ wird die Arbeit $\Delta W = F \Delta x$ geleistet. Insbesondere wirke längs der Strecke Δx_i die Kraft F_i ; hier wird dann die Arbeit $W_i = \Delta x_i F_i$ geleistet. Die Gesamtarbeit ist dementsprechend

$$W = \sum_{i=1}^n \Delta x_i F_i \quad (1.186)$$

Für $n \rightarrow \infty$ streben die Δx_i gegen Null, d.h. man kann zum Integral

$$W = \int_{x_1}^{x_2} F(\xi) d\xi \quad (1.187)$$

übergehen. Nun kann man die Geschwindigkeit als Funktion des Ortes x auffassen, $v = v(x)$. Für die Beschleunigung ergibt sich dann der Ausdruck

$$\frac{dv(t)}{dt} = \frac{dv(x)}{dx} \frac{dx(t)}{dt}.$$

Da $b(t) = dv(t)/dt$ ist kann man nun für (1.187)

$$W = \int_{x_1}^{x_2} m \frac{dv(\xi)}{d\xi} \frac{d\xi(t)}{dt} d\xi \quad (1.188)$$

schreiben. und erhält, wegen

$$\begin{aligned} \frac{dv(\xi)}{d\xi} \frac{d\xi(t)}{dt} d\xi &= v dv(\xi), \quad \text{wegen } v = \frac{d\xi}{dt}, \\ W &= \int_{x_1}^{x_2} m v dv = \frac{m}{2} v_2^2 - \frac{m}{2} v_1^2, \quad v_2 = v(x_2), \quad v_1 = v(x_1). \end{aligned} \quad (1.189)$$

Die Größe

$$T = \frac{m}{2} v^2 \quad (1.190)$$

ist die *kinetische Energie*. Der Ausdruck (1.189) bedeutet dann, dass die Arbeit W einer Veränderung der kinetischen Energie eines Körpers entspricht.

Um einen Körper P von der Position $x = 0$ zu einer Position $x \neq 0$ zu bringen, muß eine Kraft F längs des Weges von 0 bis x angewendet werden. In x wirke die Kraft $F(x)$ auf P . Hat man etwa eine Feder bis zur Position x ausgelenkt, so muß man eine Kraft $-F(x)$ aufbringen, damit sie nicht zurückschnellt. Dies kann so ausgedrückt werden, dass man sagt, dass die Feder in x ein *Potential* $U(x)$ hat, dass der Arbeit W entspricht, die geleistet werden mußte, um die Feder bis x auszulenken. Man schreibt also

$$U(x) = -W(x). \quad (1.191)$$

$U(x)$ heißt auch *potentielle Energie*. Verändert man nun W um ΔW , so ändert sich U um $-\Delta U$, und man hat

$$-\frac{dU(x)}{dx} = \frac{dW(x)}{dx} = F(x), \quad (1.192)$$

wie man aus (1.189) sieht, denn mit $x_2 = x$ erhält man (wegen $v = dx/dt$)

$$\frac{dW(x)}{dx} = 2 \frac{m}{2} v \frac{dv}{dx} = m \frac{dv}{dx} \frac{dx}{dt} = m \frac{dv}{dt} = mb(t) = F.$$

Die Gleichung (1.192) besagt dann, dass sich die Kraft als eine momentane Veränderung des Potentials auffassen läßt; dies gilt zumindest im 1-dimensionalen Fall stets. Weiter folgt

$$T(x_2) - T(x_1) = \int_{x_1}^{x_2} F(\xi) d\xi = -[U(x_2) - U(x_1)], \quad (1.193)$$

was gleichbedeutend ist mit

$$T(x_2) + U(x_2) = T(x_1) + U(x_1). \quad (1.194)$$

Diese Aussage gilt nun aber *für alle* x_1 und x_2 , so dass man auch

$$T + U = E = \text{konstant} \quad (1.195)$$

schreiben kann, – die Summe aus kinetischer und potentieller Energie ist konstant. Diese Aussage gilt allerdings nicht für alle Systeme; Systeme, für die diese Aussage gilt, heißen *konservativ*. In vielen Systemen gilt die Energieerhaltung, wie in (1.195) definiert, nicht, da Energie durch Reibungsverluste dissipiert; solche Systeme heißen dementsprechend auch *dissipative Systeme*. Dissipative Systeme sind also nicht konservativ.

Der harmonische Oszillator: Betrachtet werde ein konservatives, schwingendes System; Reibungsverluste werden also vernachlässigt. Demnach soll $T + U = E$ eine Konstante gelten. Der Auslenkung bis zur Position x entspreche nun die Potentialfunktion

$$U(x) = \frac{k}{2} x^2. \quad (1.196)$$

Da $dU(x)/dx = -F(x)$ ist /vergl. (1.192)), impliziert diese Potentialfunktion, dass die Kraft gemäß

$$F(x) = -k x \quad (1.197)$$

von x abhängt. Aus $T + U = E$ eine Konstante folgt nun

$$\frac{dT(t)}{dt} \frac{dU(t)}{dt} = 0.$$

Aus der Definition von T folgt aber

$$\frac{dT(t)}{dt} = mv \frac{dv}{dt} = mv \frac{d^2x(t)}{dt^2}, \quad \frac{dU(t)}{dt} = kx \frac{dx}{dt}.$$

Darüber hinaus muß $dT/dt = -dU/dt$ gelten, also hat man

$$mv \frac{d^2x(t)}{dt^2} = -kx \frac{dx(t)}{dt}. \quad (1.198)$$

Da nun $v = dx/dt$, folgt $md^2x(t)/dt^2 = kx(t)$, also

$$\frac{d^2x(t)}{dt^2} = -\frac{k}{m} x(t). \quad (1.199)$$

Dies bedeutet, dass die zweifache Ableitung von $x(t)$ proportional zu $-x(t)$ ist. Dies legt den Ansatz

$$x(t) = a \sin(\omega t + \phi) \quad (1.200)$$

nahe, einfach, weil die genannte Proportionalitätseigenschaft von der Sinus-Funktion erfüllt wird. In der Tat findet man

$$\frac{d^2x(t)}{dt^2} = -\omega^2 a \sin(\omega t + \phi) = -\omega^2 x(t). \quad (1.201)$$

Aus (1.199) folgt dann

$$\omega^2 = \frac{k}{m}, \quad (1.202)$$

so dass die Frequenz der Schwingung durch

$$\omega = \sqrt{k/m} \quad (1.203)$$

gegeben ist.

Die Energie des Oszillators: Wegen $T + U = E$ eine Konstante und

$$T = mv^2/2, \quad U = kx^2/2 = ka^2 \sin^2(\omega t + \phi)$$

gilt

$$T + U = \frac{mv^2}{2} + \frac{kx^2}{2} = E. \quad (1.204)$$

Weiter ist

$$v = \frac{dx(t)}{dt} = -a \cos(\omega t + \phi),$$

also

$$\frac{m}{2}a^2\omega^2 \cos^2(\omega t + \phi) + \frac{ka^2}{2} \sin^2(\omega t + \phi) = E.$$

Wegen (1.202) folgt aber

$$\frac{m}{2}a^2\omega^2 = a^2 \frac{m}{2} \frac{k}{m} = \frac{ka^2}{2},$$

und da $\cos^2(\omega t + \phi) + \sin^2(\omega t + \phi) = 1$ erhält man

$$E = \frac{k}{2} a^2. \quad (1.205)$$

Die Energie des Oszillators ist proportional zum Quadrat der Amplitude der Schwingung.

Leistung (*Power*): Leistung ist als Arbeit pro Zeiteinheit definiert. Wird an einem Körper in gleichen Zeiten gleiche Arbeit W geleistet, so ist die Leistung (der Kraft) in einem Zeitabschnitt Δt gleich $\Delta W/\Delta t$. Die Arbeit kann von der Zeit abhängen; die Leistung ist dann durch

$$L(t) = \lim_{\Delta t \rightarrow 0} \frac{\Delta W}{\Delta t} = \frac{dW(t)}{dt} \quad (1.206)$$

definiert. Schreibt man dementsprechend $L(t)dt = dW(t)$, so folgt die Beziehung

$$W(t) = \int_0^t L(\tau) d\tau. \quad (1.207)$$

Andererseits läßt sich die mittlere oder durchschnittliche Leistung definieren:

$$\bar{L} = \frac{1}{T} \int_0^T L(t)dt = \frac{1}{T} W(T). \quad (1.208)$$

Elektrizitätslehre: Die Begriffe übertragen sich auf die Elektrizitätslehre (verg. Westphal, p. 326). Ein Ladungsträger mit der Ladung Q und der Masse m durchläuft eine Spannung (Potentialdifferenz) U frei, so hat er die kinetische Energie $mv^2/2 = QU$ gewonnen. Die Strombahn (Kabel) habe einen Querschnitt q , die Geschwindigkeit, die die Ladungsträger gewonnen haben, sei v , und es seien n Ladungsträger in der Strombahn. In der Zeit dt legen sie die Strecke vdt zurück, und in einem Punkt ("Hindernis") sind so viele Ladungsträger, wie sich im Volumen $qvdt$ befinden, also $nqvdt$. Jeder hat eine kinetische Energie QU , also ist die kinetische Energie insgesamt gleich $nQqvUdt$. Aber es ist $nQv = i$ die Stromstärke in der Strombahn, – also ist die gesamte kinetische Energie gleich $Uidt$, die am Hindernis frei wird. Dies ist die *Stromarbeit* $dA = Uidt$. Sind Strom und Spannung zeitlich konstant, so hat man für die Stromarbeit in der Zeit t

$$A = Uit.$$

Die Stromleistung ist das Verhältnis von Stromarbeit und Zeit, also

$$L = dA/dt = Ui,$$

wobei U und i die momentane Spannung bzw. Stromstärke sind. Allgemein (dh U und i veränderlich) hat man

$$A = \int_0^t Uidt.$$

Allgemein gilt $U = iR$, R der Widerstand. Dann

$$dA = Uidt = i^2 Rdt, \text{ so dass } A = \int_0^t Uidt = \int_0^t i^2 Rdt.$$

Für die Leistung gilt

$$Ldt = dA = Uidt = i^2 Rdt, \text{ mithin } L = \int_0^t Ri^2 dt.$$

Diese 'Größe entspricht der insgesamt umgesetzten Energie. Diese Beziehung motiviert die folgende Definition der Signalenergie.

Signalenergie und -leistung: Es sei x ein Signal; x ist eine Funktion der Zeit. Die *Signalenergie* ist definiert durch das Integral

$$E_s = \int_{-\infty}^{\infty} |x(t)|^2 dt. \quad (1.209)$$

x kann insbesondere die Trajektorie eines stochastischen Prozesses sein, etwa $x(t) = s(t) + \xi(t)$, wobei s ein Signal und ξ die Trajektorie eines stochastischen Prozesses ξ_t , der Rauschen repräsentiert. ξ_t sei stationär. Die Energie von x wird dann eine zufällige Veränderliche sein. Die Autokorrelation von x sei durch $R(\tau)$ gegeben. Es soll nun das Leistungsspektrum von x definiert werden. Es sei

$$S_T(\omega) = \frac{1}{2T} \left| \int_{-T}^T x(t) e^{-i\omega t} dt \right|^2. \quad (1.210)$$

Dann läßt sich der folgende Satz beweisen:

Satz 1.2.14 *Es gelte*

$$\int_{-\infty}^{\infty} |\tau R(\tau)| d\tau < \infty. \quad (1.211)$$

Dann gilt

$$\lim_{T \rightarrow \infty} \mathbb{E}[S_T(\omega)] = S(\omega) = \int_{-\infty}^{\infty} R(\tau) e^{-i\omega\tau} d\tau \quad (1.212)$$

Beweis: Das Quadrat in (1.210) kann zunächst als Doppelintegral geschrieben werden:

$$S_T(\omega) = \frac{1}{2T} \int_{-T}^T x^*(u) e^{-i\omega u} du \int_{-T}^T x(v) e^{i\omega v} dv,$$

wobei x^* die konjugiert komplexe zu x sei, d.h. x darf komplex sein. Dann folgt

$$S_T(\omega) = \frac{1}{2T} \int_{-T}^T \int_{-T}^T x(u)x^*(v)e^{-i(u-v)}dudv.$$

Nun ist aber

$$\mathbb{E}[x(u)x^*(v)] = R(u-v),$$

so dass

$$\mathbb{E}[S_T(\omega)] = \frac{1}{2T} \int_{-T}^T \int_{-T}^T R(u-v)e^{-i\omega(u-v)}dudv.$$

Schreibt man nun $u-v=\tau$, so folgt⁹

$$\int_{-2T}^{2T} (2T-|\tau|)R(\tau)e^{-i\omega\tau}d\tau.$$

Dies bedeutet

$$\mathbb{E}[S_T(\omega)] = \int_{-2T}^{2T} R(\tau)e^{-i\omega\tau}d\tau - \frac{1}{2T} \int_{-2T}^{2T} |\tau|R(\tau)e^{-i\omega\tau}d\tau.$$

Das zweite Integral ist nach Voraussetzung (1.211) beschränkt, so dass wegen der Division durch $2T$

$$\frac{1}{2T} \int_{-2T}^{2T} |\tau|R(\tau)e^{-i\omega\tau}d\tau \rightarrow 0$$

folgt. Damit ist (1.212) gezeigt worden. □

Anmerkung: Die Fouriertransformierte von x sei

$$X(\omega) = \int_{-\infty}^{\infty} x(t)e^{-i\omega t}dt.$$

Nach Parsevals Satz A.3.10, Seite 170 gilt aber

$$E_x = \int_0^t |x(t)|^2 dt = \int_{-\infty}^{\infty} |X(\omega)|^2 d\omega.$$

Andererseits

$$|X(\omega)|^2 = \left| \lim_{T \rightarrow \infty} \int_{-T}^T x(t)e^{-i\omega t}dt \right|^2.$$

Nun ist $X = X_1 + iX_2$, X_1 der Real-, X_2 der Imaginärteil von X , und $|X|^2 = XX^* = X_1^2 + X_2^2$, d.h. $|X|^2$ ist das Quadrat der Amplitude, mit der ω in x eingeht, und damit gleich der Energie, die ω zur Energie von x beiträgt. $|X|^2$ und $S(\omega)$ unterscheiden sich dann, für endliches T , durch den Faktor $1/T$; deshalb ist S die durchschnittliche der mittlere Leistung bzw. Energie pro Zeiteinheit.

⁹Papoulis, p. 325

1.2.22 Linearisierung nichtlinearer Systeme

Die Nichtlinearität neuronaler, allgemein biologischer Systeme

Biologische Systeme sind im allgemeinen nichtlinear. Dies ergibt sich schon daraus, dass chemische Prozesse involviert sind, für die das Massewirkungsgesetz gilt, demzufolge die Intensität chemischer Reaktionen proportional zu den Konzentrationen der reagierenden Substanzen sind. Dieser Sachverhalt impliziert, dass bereits in die einfachsten Differentialgleichungen, die die Dynamik solcher Prozesse beschreiben, Produktterme eingehen, die die Linearität der Gleichungen zerstören. Für kleine Auslenkungen aus den Gleichgewichtslagen (den *Fixpunkten*) können die Systeme aber oft durch lineare Systeme, den *Linearisierungen* des Systems, angenähert werden. Für eine weite Klasse von Systemen ist das Verhalten der linearisierten Systeme in der Nähe der Gleichgewichtslagen qualitativ identisch ist mit dem des zugehörigen nichtlinearen Systems. Auf diese Weise liefern die linearen Annäherungen Aufschluß über die Struktur des betrachteten nichtlinearen Systems; das Verhalten des nichtlinearen Systems in der Nachbarschaft eines Fixpunktes läßt sich aus dem Verhalten des in bezug auf den Fixpunkt linearisierten Systems. Dementsprechend soll der Begriff der Linearisierung eines nichtlinearen Systems diskutiert werden.

Bei nichtlinearen Systemen kann zusätzlich zu den Fixpunkten eine Art dynamischer Gleichgewichtslage auftreten. Dies sind periodische oder oszillatorische Bewegungen. Da physiologische Befunde darauf hinweisen, dass oszillatorische Reaktionen für die Informationsverarbeitung im visuellen System möglicherweise eine zentrale Rolle spielen, soll kurz aufgezeigt werden, wie periodische Aktivitäten entstehen können.

Beispiel 1.2.15 (Aktivator-Inhibitor-Systeme) Es sei $x(t)$ die Aktivität eines neuronalen Systems zur Zeit t , $y(t)$ die eines anderen neuronalen Systems. $dx(t)/dt = \dot{x}(t)$ sei die Veränderung von x zur Zeit t , und $dy(t)/dt = \dot{y}$ sei die Veränderung von y zur Zeit t . Dazu werde der Begriff der *Aktivierungsrate* eingeführt. Diese Rate ist das Verhältnis $\dot{x}/x$ bzw. $\dot{y}/y$, also das Verhältnis der Veränderung einer Größe ($\dot{x}$ oder $\dot{y}$) zum Wert (x) bzw. (y) der Größe zu einem gegebenen Zeitpunkt; sie entspricht der Wachstumsrate in der Ökonomie oder beim Studium von Populationen. Für die Aktivierungsraten sollen die folgenden Annahmen gelten:

$$\begin{aligned}\frac{\dot{x}}{x} &= \alpha - \beta y, & \alpha > 0, & \beta > 0 \\ \frac{\dot{y}}{y} &= \gamma + \delta x, & \gamma > 0, & \delta > 0\end{aligned}\tag{1.213}$$

Für $y = 0$ soll demnach $\dot{x}/x = \alpha > 0$ gelten; dies bedeutet $\dot{x} = \alpha x$ und damit $x(t) = k_1 \exp(\alpha t)$, also exponentielles und unbeschränktes Anwachsen der Aktivierung (k_1 ist eine Konstante). Aber ein unbeschränktes Anwachsen der Aktivierung ist biologisch nicht möglich (schon weil die Spikerate von Neuronen durch die Refraktärzeit begrenzt ist). Für $y \neq 0$ soll dann der Term $-\beta y$ die nötige Dämpfung besorgen.

Für $x = 0$ soll $\dot{y}/y = -\gamma$, $\gamma > 0$ gelten, also $y(t) = k_2 e^{-\gamma t}$. Die Aktivität des zweiten neuronalen Systems geht demnach exponentiell gegen Null zurück, falls sie durch irgendeinen Prozeß von Null ausgelenkt worden ist. Der Term δx bewirkt aber, dass y mit x wächst. Für $x \neq 0$ *aktiviert* x also das y -System, und sobald $y \neq 0$ ist, übt das y -System einen *inhibierenden* Einfluß auf das x -System aus. Die Differentialgleichungen für x und y sind dann

$$\begin{aligned}\dot{x} &= x(\alpha - \beta y) \\ \dot{y} &= y(\delta x - \gamma).\end{aligned}\tag{1.214}$$

Dem Modell entspricht ein chemischer Reaktionsmechanismus (Murray (1989)): Aus einer unbeschränkt vorhandenen Substanz A plus einer Substanz X entstehen gemäß $A + X \rightarrow 2X$ zwei Teile X . Dies ist ein autokatalytischer Prozeß. Eine dritte Substanz Y liefert gemäß $X + Y \rightarrow 2Y$, ebenfalls in einem autokatalytischen Prozeß, und Y zerfällt gemäß $Y \rightarrow B$. Das Massewirkungsgesetz¹⁰ führt dann auf die Gleichungen (1.214). $X + Y \rightarrow 2Y$ wandelt X in die Substanz Y um und hemmt damit das unbeschränkte Wachsen von X . Dies ist ein Beispiel für *Substratinhibition*.

Die Gleichung (1.214) ist bekanntlich auch ein einfaches Modell für die Interaktion von Populationen: die Population A (Hasen) wächst unbeschränkt, wenn die Population B (Füchse) sie nicht daran hindert. (1.214) ist auch als Lotka-Volterra-Modell bekannt.

Nun ist $\dot{x} = dx/dt$, $\dot{y} = dy/dt$, und mithin ist $\dot{x}/\dot{y} = dx/dy$; (1.214) impliziert dann

$$\frac{dx}{dy} = \frac{x(\alpha - \beta y)}{y(\delta x - \gamma)};\tag{1.215}$$

hierüber läßt sich x als Funktion von y ausdrücken, wobei der Parameter t eliminiert wird. (1.215) liefert dann eine Darstellung der Trajektorien des Systems im Phasenraum. Aus (1.215) folgt

$$\frac{\delta x - \gamma}{x} dx = \frac{\alpha - \beta y}{y} dy$$

und die Integration liefert dann

$$\delta x - \log_e(y^\alpha x^\gamma) + \beta y = K,\tag{1.216}$$

wobei K eine Konstante ist, die sich aus den Integrationskonstanten ergibt, die bei der Integration von (1.215) resultieren. Für verschiedene Werte von K sind die Trajektorien gemäß (1.216) in angegeben. Die Trajektorien sind geschlossene Kurven; offenbar beschreiben $x(t)$ und $y(t)$ periodische Bewegungen. $\square$

Das Lotka-Volterra-System ist, wie noch deutlich werden wird, unrealistisch. Es zeigt aber exemplarisch die Entstehung von nichtlinearen Gleichungen, ist einfach und hat den Vorzug, dass geschlossene Ausdrücke für die Trajektorien sowie für $x(t)$ und $y(t)$ hergeleitet werden können. Für realistischere Modelle können solche Modelle im allgemeinen nicht hergeleitet werden, allerdings können wesentliche Eigenschaften eines solchen Systems über die Linearisierung erschlossen werden.

¹⁰Die chemische Wirkung eines Stoffes ist proportional zu seiner Konzentration.

Abbildung 1.10: Lotka-Volterra-System

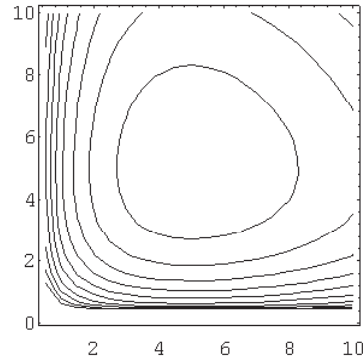
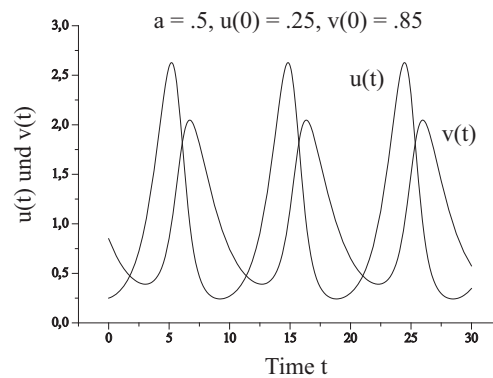


Abbildung 1.11: Lotka-Volterra-System: Trajektorien als Funktion der Zeit



Zur Linearisierung nichtlinearer Systeme

Zunächst sei daran erinnert, dass die meisten Funktionen $f(t)$ in eine Taylor-Reihe entwickelt werden können. Will man das Verhalten von f in der Nachbarschaft eines Punktes t_0 untersuchen, so gilt

$$f(t_0 + \Delta t) = f(t_0) + \Delta t f'(t_0) + \frac{\Delta t^2}{2!} f''(t_0) + \frac{\Delta t^3}{3!} f'''(t_0) + \dots$$

Für hinreichend kleines Δt gilt dann die Approximation

$$f(t_0 + \Delta t) \approx f(t_0) + \Delta t f'(t_0).$$

Eine ähnliche Entwicklung lässt sich auch für Funktionen herleiten, die von mehr als einer unabhängigen Variablen abhängen. Der Einfachheit halber sei f nun eine Funktion zweier Variablen x und y ; für $f(x_0 + x, y_0 + y)$ hat man insbesondere die Approximation

$$f(x_0 + x, y_0 + y) = f(x_0, y_0) + f_x(x_0, y_0)x + f_y(x_0, y_0)y + \dots \quad (1.217)$$

wobei f_x die partielle Ableitung von f nach x und f_y die partielle Ableitung von f nach y ist.

Gegeben sei nun ein System von zwei nichtlinearen Differentialgleichungen

$$\begin{aligned}\dot{x} &= f(x, y) \\ \dot{y} &= g(x, y)\end{aligned}\tag{1.218}$$

In Beispiel 1.2.15, (1.214), p. 74, war insbesondere $f(x, y) = x(\alpha - \beta y)$, $g(x, y) = y(\delta x - \gamma)$. Eine Linearisierung des Systems (1.218) erhält man, indem man die Funktionen f und g um einen Gleichgewichts- oder Fixpunkt (x^*, y^*) in eine Taylor-Reihe entwickelt und nur die linearen Glieder beibehält. Für einen Fixpunkt (x^*, y^*) gilt also mit $x = x^* + \Delta x$, $y = y^* + \Delta y$

$$\begin{aligned}\dot{x} &= f(x, y) = f_x(x^*, y^*)\Delta x + f_y(x^*, y^*)\Delta y \\ \dot{y} &= g(x, y) = g_x(x^*, y^*)\Delta x + g_y(x^*, y^*)\Delta y,\end{aligned}\tag{1.219}$$

da ja $f(x^*, y^*) = g(x^*, y^*) = 0$ gilt. Diese Gleichungen lassen sich in Matrixschreibweise kompakter formulieren. Dazu sei

$$A = \begin{pmatrix} f_x & f_y \\ g_x & g_y \end{pmatrix},\tag{1.220}$$

wobei die Ableitungen f_x, f_y , etc. für den Punkt (x^*, y^*) berechnet werden. A heißt auch *Jacobi-Matrix*. Statt (1.219) läßt sich dann schreiben

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \dot{x} \\ \dot{y} \end{pmatrix}\tag{1.221}$$

wobei für Δx und Δy wieder x und y geschrieben wurde. Dies ist ein lineares System mit konstanten Koeffizienten $f_x(x^*, y^*)$, $f_y(x^*, y^*)$, $\dots$. A ist die Systemmatrix. Es beschreibt das Verhalten des Systems für hinreichend kleine x und y (d.h. für die Auslenkungen Δx und Δy).

Das System (1.221) beschreibt das ursprüngliche, nichtlineare System *in einer Nachbarschaft* eines Fixpunktes in *angenäherter* Form. Genauer läßt sich das Verhalten des Systems beschreiben, wenn man die Nichtlinearität nicht ganz vernachlässigt. Dazu sei $h = (h_1, h_2)'$ ein Vektor von zwei Funktionen, die jeweils von x und y abhängen, also $h_1 = h_1(x, y)$, $h_2 = h_2(x, y)$. Dann soll gelten

$$\begin{pmatrix} \dot{x} \\ \dot{y} \end{pmatrix} = A \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} h_1(x, y) \\ h_2(x, y) \end{pmatrix}\tag{1.222}$$

h_1 und h_2 repräsentieren alle Terme der Taylor-Entwicklung (1.217) höherer Ordnung. Die lineare Approximation (1.221) ist dann vernünftig, wenn für hinreichend kleine x und y (d.h. für kleine Δx und Δy), $h_1 \rightarrow 0$ und $h_2 \rightarrow 0$. Die Art und Weise, wie diese Größen gegen Null gehen müssen, wird wie folgt spezifiziert:

Definition 1.2.12 *Es sei (1.218) ein nichtlineares System, das in der Nachbarschaft eines Fixpunktes (x^*, y^*) in der Form (1.222) geschrieben werden kann. Weiter sei $r = \sqrt{x^2 + y^2}$. Gilt nun $(h_i(x, y)/r) \rightarrow 0$ für $r \rightarrow 0$, $i = 1, 2$, so heißt das durch die Matrix A definierte lineare System (1.220) eine Linearisierung des Systems in (x^*, y^*) . Man spricht auch vom linearen Teil des Systems (1.218).*

Die Forderung $(h_i(x, y)/r) \rightarrow 0$ für $r \rightarrow 0$ spezifiziert die Art, wie die Funktionen h_i gegen Null gehen sollen: schneller, als r gegen Null geht, wenn $(x, y) \rightarrow (x^*, y^*)$. Ist diese Bedingung nicht erfüllt, so liefert das System (1.221) kein "vernünftiges" Abbild des ursprünglichen Systems. Im Übrigen ist die Definition auf Systeme mit mehr als zwei Variablen verallgemeinerbar.

Lineare Systeme mit konstanten Koeffizienten "beherrscht" man mathematisch. Die Frage ist, ob man diesen Sachverhalt nutzen kann, um nichtlineare Systeme zumindest qualitativ zu untersuchen. Dazu werden die folgenden Begriffe eingeführt:

Definition 1.2.13 Gegeben sei ein Punkt (x_0, y_0) der reellen Ebene $\mathbb{R} \times \mathbb{R}$. Eine Nachbarschaft oder Umgebung N des Punktes ist eine Teilmenge von $\mathbb{R} \times \mathbb{R}$, die Scheibe

$$\{(x, y) | \sqrt{(x - x_0)^2 + (y - y_0)^2} < r\}$$

für ein $r > 0$ enthält. Derjenige Teil eines Phasenporträts eines dynamischen Systems, der in die Nachbarschaft N fällt, heißt die Beschränkung des Phasenporträts auf N oder lokales Phasenporträt genannt.

Nichtlineare Systeme können mehr als einen Fixpunkt haben, und die lokalen Phasenporträts in der Nachbarschaft dieser Fixpunkte bestimmen nicht notwendig das globale Phasenporträt des Systems. Gleichwohl ist das Verhalten eines Systems in der Nachbarschaft der Fixpunkte von Interesse: bei instabilen Fixpunkten kann ein Übergang zu anderen Fixpunkten und damit zu einem anderen qualitativen Verhalten erfolgen. Die Frage ist, wie weit das Verhalten in der Nähe eines Fixpunktes durch das dort linearisierte System beschrieben werden kann. Ist das System in einem Fixpunkt angelangt, so findet keine Bewegung statt; Fixpunkte sind ja gerade dadurch definiert, dass in ihnen die Ableitungen der $x_i(t)$ verschwinden. Die Frage ist, wie stabil diese Ruhepunkte sind. Man betrachtet die folgenden Stabilitätsarten:

Definition 1.2.14 Ein Fixpunkt x_0 eines Systems heißt asymptotisch stabil, wenn es eine Nachbarschaft N von x_0 gibt derart, dass jede Trajektorie, die durch N geht auch gegen x_0 geht für $t \rightarrow \infty$. x_0 heißt stabil, wenn es für jede Nachbarschaft N eine in N enthaltene Nachbarschaft $\tilde{N} \subseteq N$ gibt derart, dass jede Trajektorie, die durch $\tilde{N}$ geht in N bleibt für $t \rightarrow \infty$. Ein Fixpunkt x_0 , der stabil, aber nicht asymptotisch stabil ist, heißt neutral stabil.

Jeder asymptotisch stabile Fixpunkt ist auch stabil; dies sieht man, indem man einfach $\tilde{N} = N$ wählt.

Beispiel 1.2.16 Man betrachte das durch die Dgln $\dot{x}_1 = x_2$, $\dot{x}_2 = -x_1^3$ gegebene System. Es ist stabil im Nullpunkt, aber nicht asymptotisch stabil. Man rechnet leicht nach, dass das linearisierte System durch $\dot{x}_1 = x_2$, $\dot{x}_2 = 0$ gegeben ist, und findet, dass der Fixpunkt $x_0 = 0$ nicht einfach ist. Daraus folgt, dass das linearisierte System kein lokales Phasenporträt liefert. Berücksichtigt man aber, dass $\dot{x}_1 = dx_1/dt$, $\dot{x}_2 = dx_2/dt$, so findet man $\dot{x}_2/\dot{x}_1 = dx_2/dx_1 = -x_1^3/x_2$, und diese Dgl hat Lösungen, die der Bedingung

$$\frac{1}{2}x + x = c_0 \quad \text{konstant}$$

genügen. Man ersieht aus der Bedingung, dass keine der Trajektorien den Nullpunkt $x_0 = 0$ je erreicht, was eben bedeutet, dass der Fixpunkt *nicht* asymptotisch stabil ist. Andererseits findet man für jede Umgebung N eine in N enthaltene Umgebung (Nachbarschaft) $\tilde{N}$ derart, dass jede durch $\tilde{N}$ verlaufende Trajektorie in N bleibt; $x_0 = 0$ ist also stabil, insbesondere neutral stabil. $\square$

Definition 1.2.15 *Gegeben sei ein lineares System. Das System heißt Zentrum, wenn die Dämpfungskonstante gleich Null ist.*

Die Beziehung zwischen dem Verhalten eines Systems in der Nachbarschaft eines Fixpunktes und dem des dazu korrespondierenden linearen Systems wird im folgenden Satz formuliert:

Satz 1.2.15 *Ein nichtlineares System habe einen einfachen Fixpunkt an der Stelle (x^*, y^*) . In der Nachbarschaft von (x^*, y^*) sind dann die Phasenporträts des Systems und seiner Linearisierung in (x^*, y^*) qualitativ äquivalent unter der Voraussetzung, da die Linearisierung kein Zentrum ist.*

Beweis: Knobloch und Kappel (1974), p. 207.

Die Koeffizienten des Systems, d.h. die Elemente von A , die vom speziellen Fixpunkt (x^*, y^*) abhängen, liefern ein qualitatives Bild des Verhalten des nichtlinearen Systems in der Nachbarschaft dieses Fixpunktes. Die Lösung des homogenen Systems (1.221) ist nach Satz 1.1.2, Gleichung (1.21), p. 15, durch $(x(t), y(t))' = e^{At}(x_0, y_0)'$ gegeben. Dieses Verhalten hängt von den Eigenwerten von A ab (vergl. Abschnitt 1.2.2, 22). Die Eigenwerte λ von A ergeben sich als die Wurzeln des Polynoms

$$\begin{vmatrix} f_x - \lambda & f_y \\ g_x & g_y - \lambda \end{vmatrix} = \lambda^2 - (f_x + g_y)\lambda + f_x g_y = 0. \quad (1.223)$$

Beispiel 1.2.17 Es werde noch einmal das Lotka-Volterra-System (1.214) betrachtet. Es war $f(x, y) = x(\alpha - \beta y)$, $g(x, y) = y(\delta x - \gamma)$. Das System hat offenbar zwei Fixpunkte: $(x^*, y^*) = (0, 0)$ und $(x^*, y^*) = (\gamma/\delta, \alpha/\beta)$. Allgemein findet man

$$\begin{aligned} f_x(x^*, y^*) &= \alpha - \beta y^*, & f_y(x^*, y^*) &= -\beta x^* \\ g_x(x^*, y^*) &= \delta y, & g_y(x^*, y^*) &= \delta x^* - \gamma \end{aligned}$$

Für den Fixpunkt $(x^*, y^*) = (0, 0)$ erhält man die Eigenwerte

$$\begin{vmatrix} \alpha - \lambda & 0 \\ 0 & -c - \lambda \end{vmatrix} = -(\alpha - \lambda)(\gamma + \lambda) = 0.$$

Diese Gleichung ist erfüllt, wenn entweder $\alpha - \lambda = 0$ oder $\gamma + \lambda = 0$, also gibt es stets zwei Lösungen für λ , $\lambda_1 = \alpha$, $\lambda_2 = -\gamma$, denn $\alpha > 0$, $\gamma > 0$ nach Voraussetzung. Nach (1.39), p. 22, ist aber $e^{At} = T e^{\Lambda t} T^{-1}$, $\Lambda = \text{diag}(\lambda_1, \lambda_2)$, und ein positiver Eigenwert bedeutet, dass das entsprechende Element der Impulsantwort unbeschränkt wächst.

Damit ist gezeigt, dass der Fixpunkt $(0,0)$ des Lotka-Volterra-Systems instabil ist, die kleinste Störung treibt das System von diesem Gleichgewichtspunkt fort.

Für den zweiten Fixpunkt ergibt sich die Determinante

$$A = \begin{vmatrix} -\lambda & -\beta\gamma/\delta \\ \alpha\delta/\beta & -\lambda \end{vmatrix} = \lambda^2 + \alpha\gamma = 0$$

und man erhält die beiden Eigenwerte

$$\lambda_{1,2} = \pm \frac{1}{2} \sqrt{-4\alpha\gamma} = \pm \frac{i}{2} \sqrt{4\alpha\gamma}.$$

Die Eigenwerte sind rein imaginär, das linearisierte System ist ein *Zentrum*. Dies bedeutet, dass jede Auslenkung des Systems aus dem Fixpunkt $(\gamma/\delta, \alpha/\beta)$ zu einer periodischen, *nicht abklingenden* Antwort führt. Dieser Sachverhalt entspricht den in (1.216) gefundenen geschlossenen Orbits im Phasenraum. Jede kleine Störung, die zu einer Abweichung von einem solchen Orbit führt, veranlaßt das System sofort dazu, zu einem neuen Orbit überzugehen und bei diesem zu bleiben, so lange es nicht weiter gestört wird. Dies ist eine der Eigenschaften, die das System unrealistisch sein lassen. Man sagt, das System sei *strukturell instabil*. $\square$

Für ein stabiles, lineares System mit einem wohldefinierten Fixpunkt gibt es qualitativ gesehen nur eine Reaktion auf eine Auslenkung: die Trajektorie bewegt sich auf den Fixpunkt zu. Ist er erreicht, so verbleibt das System in Ruhe. In nichtlinearen Systemen kann ein anderes Verhalten auftreten. Die Trajektorie kann sich u. U. an eine geschlossene Trajektorie, die kein Punkt ist, annähern. Wird diese Trajektorie erreicht, so wird die Bewegung des Systems durch diese geschlossene Trajektorie beschrieben. Die Geschlossenheit dieser Trajektorie bedeutet, dass eine bestimmte Zustandsmenge zyklisch durchlaufen wird. Dementsprechend hat man die folgende

Definition 1.2.16 *Eine geschlossene Trajektorie in einem Phasenporträt heißt Grenzyklus genau dann, wenn sie von allen anderen geschlossenen Trajektorien getrennt ist.*

Aus den vorangegangenen Bemerkungen ist die Wahl des Ausdrucks *Grenzyklus* klar. Ein Grenzyklus ist dann gegeben, wenn es einen Tubus gibt, in dem die betrachtete geschlossene Trajektorie liegt und in dem es keine andere geschlossene Trajektorie gibt. Ein Beispiel verdeutlicht den Begriff:

Beispiel 1.2.18 Gegeben sei das System¹¹

$$\begin{aligned} \dot{x}_1 &= -x_2 + x_1(1-r) \\ \dot{x}_2 &= x_1 + x_2(1-r) \end{aligned} \tag{1.224}$$

mit $r = \sqrt{x_1^2 + x_2^2}$. Das System hat einen durch die Gleichung $x_1^2 + x_2^2 = 1$ gegebenen Grenzyklus. Denn setzt man $x_1 = r \cos \theta$, $x_2 = r \sin \theta$, so geht (1.224) über in

¹¹Arrowsmith und Place (1982), p. 107

$\dot{r} = r(1 - r)$, $\dot{\theta} = 1$. Man überprüft leicht, dass $r(t) \equiv 1$, $\theta(t) = t$ eine Lösung ist, die einer geschlossenen Trajektorie entspricht, die der Bedingung $x + x = 1$ genügt und diesen Kreis entgegen dem Uhrzeigersinn mit konstanter Winkelgeschwindigkeit $\dot{\theta} = 1$ durchläuft. Für $0 < r < 1$ ist $\dot{r} > 0$ und die Trajektorien spiralen nach außen gegen den Wert $r = 1$. Für $r > 1$ ist $\dot{r} < 0$ und die Trajektorien spiralen nach innen mit steigendem t . $\square$

Es gibt drei Arten von Grenzzyklen:

- Der stabile Grenzzyklus. Für $t \rightarrow \infty$ spiralen die Trajektorien von außen sowie von innen gegen den Grenzzyklus.
- Den instabilen Grenzzyklus. Für $t \rightarrow \infty$ bewegen sich hier die Trajektorien vom Grenzzyklus weg in beide Richtungen.
- Der semistabile Grenzzyklus. Hier bewegen sich die Trajektorien von der einen Seite auf den Grenzzyklus zu und bewegen sich auf der anderen Seite von ihm fort.

Grenzzyklen müssen natürlich nicht kreisförmig sein und werden dementsprechend nicht notwendig durch Übergang zu Polarkoordinaten erfaßt. Kriterien für das Vorhandensein von Grenzzyklen werden in der Poincaré — Bendixson Theorie behandelt, auf die an dieser Stelle nicht weiter eingegangen werden kann.

Es soll noch kurz betrachtet werden, was geschieht, wenn sich die Parameter dynamischen Systems verändern.

Dazu werde zunächst das einfache System $\dot{x} = ax$ betrachtet. Offenbar ist $x = be^{at}$, $b \in \mathbb{R}$, eine Lösung. Für $a > 0$ wächst x unbegrenzt, d.h. es bewegt sich vom Nullpunkt fort: der Nullpunkt ist ein *Repellor*. Für $a < 0$ bewegt es sich auf den Nullpunkt zu, der Nullpunkt ist dann ein *Attraktor*. Für $a < 0$ und für $a > 0$ erhält man also verschiedene Phasenporträts. Anders ausgedrückt heißt dies, dass man beim Übergang von $a < 0$ zu $a > 0$ ein anderes qualitatives Verhalten des Systems erhält. Der Punkt, an dem das Verhalten des Systems sich qualitativ verändert, ist $a = 0$. Dieser Punkt heißt *Bifurkationspunkt*; das System zeigt in $a = 0$ eine *Bifurkation*.

Das System $\dot{x}_1 = \mu x$, $\dot{x}_2 = -x_2$ zeigt ebenfalls für $\mu = 0$ eine Bifurkation, s.

Fig. 2.9.5.

Das folgende Theorem ist für die qualitative Betrachtung nichtlinearer Systeme von großer Bedeutung:

Satz 1.2.16 *Betrachtet werde das von dem Parameter μ abhängige System*

$$\dot{x}_1 = f_1(x_1, x_2, \mu), \quad \dot{x}_2 = f_2(x_1, x_2, \mu). \quad (1.225)$$

Das System habe den Fixpunkt $x_0 = 0$ für alle Werte des Parameters μ , und für $\mu = \mu_0$ seien die Eigenwerte $\lambda_1(\mu)$, $\lambda_2(\mu)$ des linearisierten Systems rein imaginär.

Für $\mu \neq \mu_0$ gelte für den Realteil $\alpha = \operatorname{Re}(\lambda_1(\mu)) = \operatorname{Re}(\lambda_2(\mu))$ die Bedingung

$$\frac{d\alpha}{d\mu}\bigg|_{\mu=\mu_0} > 0 \quad (1.226)$$

und $x_0 = 0$ sei asymptotisch stabil. Dann gilt

- a. $\mu = \mu_0$ ist ein Bifurkationspunkt des Systems.
- b. Es existiert ein $\mu_1 < \mu$ derart, dass für alle $\mu \in (\mu_1, \mu)$ der Punkt $x_0 = 0$ ein stabiler Fokus ist.
- c. Es existiert ein $\mu_2 > \mu$ derart, dass für alle $\mu \in (\mu_0, \mu_2)$ der Punkt $x_0 = 0$ ein instabiler Fokus ist, der durch einen stabilen Grenzzyklus umgeben ist. Die Größe des Grenzzyklus wächst mit μ .

Beweis: Ein Beweis kann hier nicht gegeben werden; vergl. Marsden et al. (1976).

Bifurkationen der in Satz 1.2.16 beschriebenen Art heißen *Hopf-Bifurkation*. Hopf-Bifurkationen spielen eine wichtige Rolle bei der Betrachtung komplexer Systeme (vergl. Nicolis und Prigogine (1987)).

Fig. 2.9.5 Bifurkation von $\dot{x}_1 = \mu x_1, \dot{x}_2 = -x_2$, (a) $\mu < 0$, (b) $\mu = 0$, (c) $\mu > 0$

Hier Kapitel über Entdeckensmechanismen

Kapitel 2

Die allgemeine Theorie der Signalentdeckung

Die allgemeine Theorie der Signalentdeckung soll hier nicht in aller Allgemeinheit dargestellt werden (vergl. Whalen (1971), Macmillan and Creelman (2005)), sondern nur so weit, wie es nötig ist.

Der Begriff der Schwellenintensität oder des Schwellenkontrasts c bezieht sich auf eine Wahrscheinlichkeit des Entdeckens: gegeben ein bestimmter Wert von c , soll die Wahrscheinlichkeit des Entdeckens gleich $P(c)$ sein. Dabei wird c hinreichend klein gewählt, so dass $P(c) < 1$. Für einen derartigen c -Wert wird der Stimulus also in einigen Versuchsdurchgängen wahrgenommen – entdeckt –, und in einigen nicht. Dies bedeutet, dass die Wahrnehmung einerseits und der Stimulus mit gegebenem c -Wert nicht deterministisch, sondern stochastisch gekoppelt sind. Als Hintergrund für diese Art der Kopplung kann man statistische Fluktuationen in der Stimuluspräsentation (Photonen sind Poisson-verteilt) und Fluktuationen in der Stimulusverarbeitung annehmen. Bei den meisten Experimenten werden die Stimuli vor einem Hintergrund mit Luminanz $L_0 > 0$ dargeboten, d.h. die Versuchsperson ist nicht vollständig dunkeladaptiert und die Stimulusintensität ist so hoch, dass zwar Photonendiffluktuationen existieren, aber relativ zur durchschnittlichen Anzahl emittierter Photonen vernachlässigbar sind. Der Fokus der Modellierung des Entdeckensprozesses wird also auf den Fluktuationen in den Signalverarbeitenden Mechanismen liegen. Eine Möglichkeit besteht darin, Fluktuationen in den "höheren", entscheidenden Zentren zu suchen; hier kämen zufällige Schwankungen der Aufmerksamkeit in Frage. Eine andere Möglichkeit besteht darin, zufällige Effekte in der primären Signalverarbeitung anzunehmen. In der Tat ist die Aktivierung individueller Neurone als stochastischer Prozess zu beschreiben. In angenäherter Form sind die Folgen von Aktionspotentialen (Spikes) in guter Näherung als Poisson-Folgen zu beschreiben (wegen der Refraktärzeit nach erfolgter Spike-Produktion kann die Spike-Folge nicht einem echten Poisson-Prozess entsprechen). Am Entscheidungszentrum wird die Aktivität dementsprechend ebenfalls durch einen stochastischen Prozess zu beschreiben sein; die Aufgabe der Versuchsperson (V_p) besteht darin, zu entscheiden, ob die beobachtete Aktivität als durch einen Stimulus generiert oder nicht durch einen Stimulus

Tabelle 2.1: Entscheidung und Wirklichkeit (*state of nature*; "ja" steht für "Stimulus wahrgenommen", "nein" für "Stimulus nicht wahrgenommen".)

	state of nature	
Entscheidung	Stimulus (σ)	kein Stimulus ($\bar{\sigma}$)
ja	$1 - \beta$	α
nein	β	$1 - \alpha$

generiert anzusehen ist. Damit wird man sofort auf das übliche, in Tabelle 2.1 angegebene entscheidungstheoretische Schema geführt. Darin ist α die Wahrscheinlichkeit, einen Reiz wahrzunehmen, obwohl kein Reiz präsentiert wurde (falscher Alarm), und β ist die Wahrscheinlichkeit, keinen Reiz wahrzunehmen, obwohl einer präsentiert wurde. Ist D das zufällige Ereignis, dass die Vp meint, einen Reiz wahrgenommen zu haben, und $\bar{D}$ das dazu komplementäre Ereignis, keinen Reiz wahrgenommen zu haben, so hat man also

$$\alpha = P(D|\bar{\sigma}) \quad (2.1)$$

$$\beta = P(\bar{D}|\sigma) \quad (2.2)$$

Die Werte von α und β hängen einerseits von der Stimulusintensität, andererseits vom Kriterium der Vp ab, nach dem die Vp urteilt, ob die Aktivität in einem Versuchsdurchgang indikativ für einen Reiz ist oder nicht.

Es wird im Folgenden davon ausgegangen, dass die neuronale Aktivität, die zum Entdecken eines Stimulus führt, als stochastischer Prozess charakterisiert werden muß, wobei die mittlere Aktivität durch eine deterministische Funktion – also eine Funktion, in deren Definition keine Terme eingehen, die zufällige Größen repräsentieren – beschrieben werden kann. Im Allgemeinen wird dann angenommen, dass diese Funktion durch ein deterministisches dynamisches System spezifiziert werden kann. Grundbegriffe der Theorie stochastischer Prozesse werden in Abschnitt A.9, Seite 185 gegeben. Demnach werden im Folgenden mit η_t bzw. ξ_t stochastische Prozesse bezeichnet; dies sind Familien von Funktionen der Zeit η oder $\eta(\cdot)$ bzw. ξ oder $\xi(\cdot)$. η und ξ sind die *Trajektorien* des jeweiligen Prozesses. Gelegentlich wird auch $\eta(t)$ oder $\xi(t)$ geschrieben, um anzuzeigen, dass eine bestimmte Funktion aus η_t oder ξ_t gemeint ist. Andererseits sind $\eta(t)$ und $\xi(t)$ natürlich die Werte der Funktionen η und ξ .

Es sei nun η_t der die relevante neuronale Aktivität repräsentierende Prozess,

wenn ein Stimulus gezeigt wurde, und ξ_t sei dieser Prozess, wenn kein Stimulus gezeigt wurde. Es wird angenommen, dass

$$\eta_t = \xi_t + g, \quad (2.3)$$

wobei g eine deterministische Funktion der Zeit ist, die die durchschnittliche Aktivität abbildet. Dies bedeutet

$$\mathbb{E}(\eta_t) = g, \quad \mathbb{E}(\xi_t) = 0. \quad (2.4)$$

Mit $\mathbb{E}$ ist der Erwartungswert gemeint. Demnach ist $\xi_t = \eta_t - g$, ξ_t ist also zunächst einfach die Abweichung von der mittleren Aktivität. Nimmt man zusätzlich an, dass ξ_t stationär ist, so wird impliziert, dass ξ_t unabhängig von g ist. Diese Annahme ist nicht trivial. Betrachtet man etwa für η_t einen Poisson-Prozess, so gilt diese Annahme nicht mehr. η sei die Funktion, die die Aktivität während eines Versuchsdurchganges repräsentiert. Die Vpn steht vor der Aufgabe, zu entscheiden, ob $\eta \in \eta_t$ oder $\eta \in \xi_t$ gilt, – wird also davon ausgegangen, dass $\eta_t \cap \xi_t \neq \emptyset$ gilt. Formal lassen sich die Hypothesen, zwischen denen sich die Versuchsperson entscheiden muß, durch

$$\begin{aligned} H_0 : & \quad \eta_t = \xi_t \\ H_1 : & \quad \eta_t = \xi_t + g, \end{aligned} \quad 0 \leq t \leq T \quad (2.5)$$

charakterisieren. Die bedingten Wahrscheinlichkeiten $P(D|\sigma)$, $P(D|\bar{\sigma})$ etc. müssen also in Bezug auf die Trajektorie η getroffen werden. Man kann annehmen, dass die Entscheidung nach Maßgabe des Wertes einer zufälligen Veränderlichen X vorgenommen wird; für $X > \eta_0$ wird ein Stimulus wahrgenommen, für $X \leq S$ nicht. Die Wahrscheinlichkeit $\psi(c)$, den Stimulus zu entdecken, wenn die Stimulusintensität gleich c ist – läßt sich über die Verteilungsfunktion von X definieren:

$$\psi(c) = 1 - P(X \leq \eta_0). \quad (2.6)$$

Die nächste Frage ist nun, wie die zufällige Veränderliche X und damit die Verteilungsfunktion $P(X \leq x)$ definiert ist.

2.1 Entdeckenswahrscheinlichkeiten als bedingte Wahrscheinlichkeiten

Es seien $P(H_0|\eta)$ und $P(H_1|\eta)$ die bedingten Wahrscheinlichkeiten für die beiden Hypothesen H_0 und H_1 , gegeben die durch η repräsentierte Aktivität. Man kann dann

$$P(H_0|\eta) = \frac{P(\eta|H_0)p(H_0)}{P(\eta)}, \quad P(H_1|\eta) = \frac{P(\eta|H_1)P(H_1)}{P(\eta)}$$

schreiben. Dann ist

$$\frac{P(H_1|\eta)}{P(H_0|\eta)} = \frac{P(\eta|H_1)}{P(\eta|H_0)} \frac{P(H_1)}{P(H_0)} \quad (2.7)$$

Die Wahrscheinlichkeiten $P(H_0)$ und $P(H_1)$ können als *a priori*-Wahrscheinlichkeiten für einen Versuchsdurchgang mit oder ohne Stimulusdarbietung betrachtet werden. Der Quotient

$$\lambda(\eta) = \frac{P(\eta|H_1)}{P(\eta|H_0)} \quad (2.8)$$

ist hierin der *Likelihood-Quotient*. Die Versuchsperson könnte nun nach Maßgabe des Wertes des Quotienten $P(H_1|\eta)/P(H_0|\eta)$ entscheiden, also nach

$$\frac{P(H_1|\eta)}{P(H_0|\eta)} = \lambda(\eta) \frac{P(H_1)}{P(H_0)} = \lambda(\eta) \frac{P(H_1)}{1 - P(H_1)}. \quad (2.9)$$

Dazu muß angenommen werden, dass sie nicht nur den Wert von $\lambda(\eta)$ evaluieren kann, sondern auch noch den Quotienten $P(H_1)/P(H_0)$ mit in ihr Urteil eingehen lassen kann. Alternativ dazu kann man annehmen, dass sie sich nur nach dem Wert von $\lambda(\eta)$ richtet. $\lambda(\eta)$ ist eine zufällige Veränderliche und die Vp entscheidet, es sei ein Stimulus präsentiert worden, wenn $\lambda(\eta) > \eta_0$; es wird also angenommen, dass $X = \lambda(\eta)$. Der Wert kann implizit durch Variation der pay-off-Bedingungen vom Experimentator manipuliert werden, d.h. es wird angenommen, dass die Vp die Möglichkeit hat, den Wert von η_0 implizit zu setzen.

Es muß noch geklärt werden, in welcher Beziehung die Wahrscheinlichkeiten $P(H_i|\eta)$ und $P(\eta|H_i)$, $i = 0, 1$, zu den Wahrscheinlichkeiten $P(D|\sigma)$, $P(\bar{D}|\sigma)$, $P(D|\bar{\sigma})$ und $P(\bar{D}|\bar{\sigma})$ stehen. Denn zunächst einmal können die $P(\eta|H_i)$ und $P(H_i)$, $i = 0, 1$, als *objektive* Wahrscheinlichkeiten gesehen werden. Damit wären dann auch die $P(H_i|\eta)$ objektive Größen. Setzt man nun

$$P(\eta|H_i) = \begin{cases} P(\eta|\sigma), & i = 1, \text{ d.h. Stimulus präsentiert} \\ P(\eta|\bar{\sigma}), & i = 0, \text{ d.h. kein Stimulus präsentiert} \end{cases}, \quad (2.10)$$

so impliziert man, dass die jeweilige objektive Wahrscheinlichkeit von η gleich der entsprechenden subjektiven Wahrscheinlichkeit ist, mit der die Versuchsperson die neuronale Aktivität mit der Möglichkeit einer Reizdarbietung in Beziehung setzt. Ob die Beziehung (2.10) tatsächlich gilt, ist eine Frage, die empirisch entschieden werden muß. Andererseits kann man auch einfach die beobachteten bedingten Wahrscheinlichkeiten $P(D|\sigma)$ etc. betrachten. Die *Receiver-Operator-Characteristics* (ROC), d.h. die Plots von $P(D|\eta_0)$ versus $P(D|\bar{\eta}_0)$, für konstante Stimulusintensität zeigen, dass eine Manipulation des Schwellenwertes η_0 tatsächlich möglich ist. Unter der üblichen Annahme, dass ξ_t Gaußschen Weißen Rauschens ist, läßt sich zeigen, dass λ Gauss-verteilt ist. Damit hat man eine zufällige Veränderliche derart, dass

$$P(D|\eta_0) = P(\lambda > \lambda_0|\eta_0), \quad (2.11)$$

d.h. die Wahrscheinlichkeit des Entdeckens läßt sich über die Verteilungsfunktion von $X = \lambda$ definieren.

2.2 Der Likelihood-Quotient für Weißes Rauschen

Eine Möglichkeit, die Wahrscheinlichkeiten $P(\eta|s)$, $P(\eta|\bar{\sigma})$ und damit $\lambda(\eta)$ zu schätzen besteht darin, das Signal $\eta(t)$ an m Stellen abzutasten; man erhält die Werte

$$\eta(t_k) = \begin{cases} s(t_k) + \xi(t_k), & \text{Stimulus} \\ \xi(t_k), & \text{kein Stimulus} \end{cases}, k = 1, \dots, m \quad (2.12)$$

Die Wahrscheinlichkeit $P(\eta|s)$ läßt sich nun als Grenzwert

$$P(\eta|s) = \lim_{m \rightarrow \infty} P[\eta(t_1) \cap \eta(t_2) \cap \dots \cap \eta(t_m)], \quad \eta(t_k) = \sigma(t_k) + \xi(t_k), \quad 1 \leq k \leq m \quad (2.13)$$

ansetzen; der Ausdruck für $P(\eta|\bar{\sigma})$ ist analog. Könnte man die zufälligen Variablen $\eta(t_k)$ als stochastisch unabhängig annehmen (wie es Forscher in der visuellen Psychophysik es gelegentlich tun), wäre man schnell fertig. Es wird aber gezeigt werden, dass diese Annahme zu inkonsistenten Modellen des Entdeckens führt. Man kann aber die – immer noch vereinfachende – Annahme machen, dass die Autokorrelation des Rauschens (das durch die Trajektorien $\xi(t)$ repräsentiert wird – durch

$$R(\tau) = \frac{N_0 \Omega}{2\pi} \frac{\sin(\Omega\tau)}{\Omega\tau} \quad (2.14)$$

gegeben ist, wobei die implizite Annahme ist, dass das Rauschen ein stationärer Prozess ist. Diese auf den ersten Blick willkürlich anmutende Annahme über $R(\tau)$ wird motiviert durch eine vereinfachende Annahme über die Spektraldichte des Rauschens. Die Spektraldichte ist als Fouriertransformierte der Autokorrelationsfunktion definiert (vergl. Definition A.9.6, Seite 188):

$$F(\omega) = \int_{-\infty}^{\infty} R(\tau) e^{-i\omega\tau} d\tau, \quad (2.15)$$

und natürlich gilt dann auch

$$R(\tau) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega\tau} d\omega \quad (2.16)$$

Man kann zwei Fälle betrachten: (i) das weiße Rauschen:

$$F(\omega) = K, \quad -\infty < \omega < \infty, \quad (2.17)$$

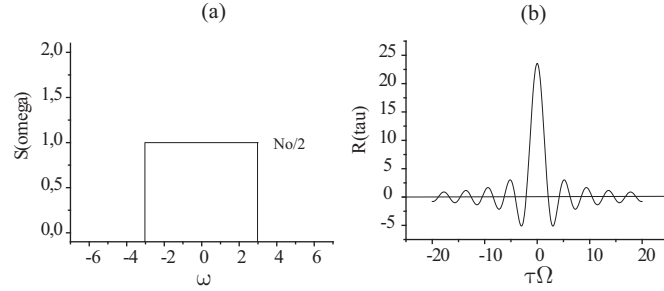
und (ii) das bandbegrenzte Rauschen

$$F(\omega) = K = N_0/2, \quad |\omega| < \Omega \quad (2.18)$$

Für das weiße Rauschen (??) liefert (vergl. Abschnitt A.4.3, Seite 173) die Beziehung (2.16)

$$R(\tau) = \frac{K}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega\tau} d\omega = K\delta(\tau). \quad (2.19)$$

Abbildung 2.1: Spektraldichte (a) und korrespondierende Autokorrelationsfunktion eines stationären Gauss-Prozesses (b). $R(\tau) = 0$ für $\tau = k\pi/\Omega$, $k = \pm 1, \pm 2, \dots$



Dies heißt $R(\tau) = 0$ für alle $\tau \neq 0$, und für einen Gauß-Prozess (die $\xi(t)$, also die Werte von ξ zum Zeitpunkt t , sind Gauß-verteilt) sind $\xi(t)$ und $\xi(t')$ für *alle* t, t' mit $t \neq t'$ stochastisch unabhängig. Das Rauschen ist dann *Weißes Rauschen*. Das Rauschen heißt 'weiß', weil wegen $S(\omega) = K$ wie beim weißen Licht alle Frequenzen mit konstanter Amplitude benötigt werden, um eine beliebige Trajektorie als Superposition von Sinus-Funktionen darzustellen. Dies wiederum bedeutet, dass der Prozess beliebig schnell fluktuiert. Es bedeutet darüber hinaus, dass die Energie

$$E_\xi = \int_{-\infty}^{\infty} |F(\omega)|^2 d\omega = \int_{-\infty}^{\infty} K^2 d\omega = \infty \quad (2.20)$$

des Prozesses unendlich ist. Das Weiße Rauschen ist also allenfalls eine Idealisierung, die physikalisch nicht realisierbar ist. Weitere unplausible Implikationen einer allzu freizügigen Annahme des Weißen Rauschens werden im Zusammenhang mit der Definition von psychometrischen Funktionen angegeben.

Für das bandbegrenzte weiße Rauschen erhält man aus (2.16) eben die Autokorrelationsfunktion (2.14), und für die Energie E_ξ des Prozesses ξ_t erhält man gemäß (2.20)

$$E_\xi = \int_{-\infty}^{\infty} |F(\omega)|^2 d\omega = \frac{N_0}{2} \int_{-\Omega}^{\Omega} d\omega = N_0\Omega < \infty. \quad (2.21)$$

Der erste Nulldurchgang von $R(\tau)$ erfolgt zum Zeitpunkt $\tau = \pi/\Omega$. Da die Nulldurchgänge gleichabständig sind, sind die zufälligen Veränderlichen für $\tau_k = k\pi/\Omega$ stochastisch unabhängig (dies folgt wegen der zusätzlichen Annahme, dass ξ_t ein Gauss-Prozess ist). Im Zeitintervall $(0, T)$ kann man

$$m = \frac{T}{\Delta t} = \frac{\Omega T}{\pi} \quad (2.22)$$

unabhängige Werte $\eta(t_j) = \eta_j$, $j = 1, \dots, m$ messen. Der Erwartungswert und die Varianz von η_j sind

$$\mathbb{E}(\eta_j) = \begin{cases} s(t_j), & \text{Stimulus präsentiert} \\ 0, & \text{kein Stimulus präsentiert} \end{cases} \quad (2.23)$$

und

$$\mathbb{V}(\eta_j) = \mathbb{V}(\xi(t_j)) = \sigma^2 = \frac{N_0 \Omega}{2\pi}. \quad (2.24)$$

Für die Likelihood-Funktionen beginnt man mit dem Ansatz

$$P(\eta|s) = \left(\frac{1}{2\pi\sigma^2} \right)^{m/2} \exp \left[- \sum_{j=1}^m \frac{(\eta_j - s(t_j))^2}{2\sigma^2} \right] \quad (2.25)$$

und

$$P(\eta|\bar{s}) = \left(\frac{1}{2\pi\sigma^2} \right)^{m/2} \exp \left[- \sum_{j=1}^m \frac{\eta_j^2}{2\sigma^2} \right] \quad (2.26)$$

und erhält den Likelihood-Quotienten

$$\lambda(\eta) = \frac{P(\eta|s)}{P(\eta|\bar{s})} = \frac{\exp \left[- \sum_{j=1}^m \frac{(\eta_j - s(t_j))^2}{2\sigma^2} \right]}{\exp \left[- \sum_{j=1}^m \frac{\eta_j^2}{2\sigma^2} \right]}. \quad (2.27)$$

Dieser Ausdruck kann in der Form

$$\lambda(\eta) = \exp \left[\sum_{j=1}^m \frac{\eta_j^2}{2\sigma^2} - \sum_{j=1}^m \frac{(\eta_j - s_j)^2}{2\sigma^2} \right], \quad s_j = s(t_j)$$

geschrieben werden, und wenn die Entscheidung, dass ein Stimulus präsentiert wurde gefällt werden soll, wenn $\lambda(\eta) > \lambda_0$. Nun ist

$$\begin{aligned} \log \lambda(\eta) &= \sum_{j=1}^m \frac{\eta_j^2}{2\sigma^2} - \sum_{j=1}^m \frac{(\eta_j - s_j)^2}{2\sigma^2} = \frac{1}{2\sigma^2} \left[\sum_{j=1}^m \eta_j^2 - \sum_{j=1}^m (\eta_j - s_j)^2 \right] \\ &= \frac{1}{2\sigma^2} \left[\sum_{j=1}^m \eta_j^2 - \sum_{j=1}^m \eta_j^2 - \sum_{j=1}^m s_j^2 + 2 \sum_{j=1}^m \eta_j s_j \right] \\ &= \frac{1}{\sigma^2} \sum_{j=1}^m \eta_j s_j - \frac{1}{2\sigma^2} \sum_{j=1}^m s_j^2 \end{aligned} \quad (2.28)$$

Die Summen lassen sich durch einen Grenzwertprozess in Integrale verwandeln. Denn aus (2.22) und (2.24) folgt

$$\frac{1}{\sigma^2} = \frac{2\pi}{\Omega N_0} = \frac{2}{N_0} \Delta t, \quad (2.29)$$

so dass (2.28) zu

$$\log \lambda(\eta) = \frac{2}{N_0} \sum_{j=1}^m \eta_j s_j \Delta t - \frac{2}{N_0} \frac{1}{2} \sum_{j=1}^m s_j^2 \Delta t \quad (2.30)$$

wird. Für $m \rightarrow \infty$ und $\Delta t \rightarrow 0$ derart, dass $T = m\Delta t = \text{konstant}$ kann das $\sum$ -Zeichen durch ein $\int$ -Zeichen und Δt durch dt ersetzt werden und man erhält die Gleichung

$$\log \lambda(\eta) = \frac{2}{N_0} \int_0^T \eta(t)s(t)dt - \frac{1}{N_0} \int_0^T s^2(t)dt. \quad (2.31)$$

Hierin ist

$$\int_0^T s^2(t)dt = \int_0^T |s(t)|^2 dt = E_s$$

die Energie des Signals im Intervall $[0, T]$. Das Integral

$$R_T = \int_0^T \eta(t)s(t)dt \quad (2.32)$$

kann als Korrelation zwischen der Aktivität und dem Signal aufgefasst werden. R_T ist eine zufällige Veränderliche, da η eine Trajektorie eines stochastischen Prozesses η_t ist.

Der Übergang $m \rightarrow \infty$ impliziert $\Delta t \rightarrow 0$, so dass nach (2.29) $\Omega \rightarrow \infty$ folgt. Damit ist aber das Spektrum $S(\omega)$ des Rauschens nicht mehr auf ein endliches Intervall begrenzt, und $R(\tau) \rightarrow (N_0/2)\delta(\tau)$, d.h. der Grenzübergang impliziert einen Übergang zum Weißen Rauschen. In (2.30) tritt die Summe $\sum_{j=1}^m \eta_j s_j \Delta t$ auf, die in (2.31) durch das Integral $\int_0^T \eta(t)s(t)dt$ ersetzt wird. Die η_j sind unabhängige Gauß-Variable, und die ("gewichtete") Summe von unabhängigen Gauß-Variablen ist wiederum Gauß-verteilt. Dieses Merkmal vererbt sich, wenn weißes Rauschen angenommen wird, auf das Integral. Damit folgt, dass die zufällige Veränderliche R_T normalverteilt ist.

$X = \log \lambda(\eta)$ ist eine zufällige Veränderliche, von deren Wert die Entscheidung über die Darbietung eines Stimulus abhängt; dies erklärt den gelegentlich benützten Ausdruck *Entscheidungsvariable*. Man kann demnach

$$X = \log \lambda(\eta) = R_T - \frac{E_s}{N_0} \quad (2.33)$$

schreiben. Hierin ist E_s/N_0 eine Konstante. Der Erwartungswert und die Varianz sind dann durch

$$\mathbb{E}(X) = \mathbb{E}(R_T) - \frac{E_s}{N_0}, \quad \mathbb{V}(X) = \mathbb{E}(R_T^2) - \mathbb{E}^2(R_T) \quad (2.34)$$

gegeben. Unter H_1 ist $\eta(t) = \xi(t) + g(t)$ und $\mathbb{E}(\eta(t)) = g(t)$, also $\mathbb{E}(\xi(t)) = 0$ für alle t . Dann ist

$$\int_0^T \eta(t)s(t)dt = \int_0^T \xi(t)s(t)dt + \int_0^T g(t)s(t)dt, \quad (2.35)$$

so dass

$$\mathbb{E}(R_T|H_1) = \mu_1 = \int_0^T g(t)s(t)dt, \quad (2.36)$$

denn

$$\mathbb{E} \left[\int_0^T \xi(t)s(t)dt \right] = \int_0^T \mathbb{E}[\xi(t)]s(t)dt = 0, \quad \text{wegen } \mathbb{E}[\xi(t)] \equiv 0,$$

woraus sofort auch

$$\mathbb{E}(R_T|H_0) = \mu_0 = 0 \quad (2.37)$$

folgt, da unter H_0 $\eta(t) = \xi(t)$.

Varianz: noch offen!!!

Der Erwartungswert von X unter H_1 ist dann

$$\mathbb{E}(X|H_1) = \mu_1 - \frac{E_s}{N_0}, \quad (2.38)$$

und unter H_0

$$\mathbb{E}(X|H_0) = -\frac{E_s}{N_0}. \quad (2.39)$$

Die Differenz

$$d' = \frac{\mu_1 - \mu_0}{\sigma_x} = \frac{1}{\sigma_x} \frac{E_s}{N_0}, \quad (2.40)$$

wobei $\sigma_x = \sqrt{\mathbb{V}(X)}$, heißt *Sensitivitätsmaß*, und der Quotient E_s/N_0 ist das Signal-Rausch-Verhältnis. Allgemein ist das Signal-Rausch-Verhältnis – abgekürzt SRV, oder SNR, wegen des englischen Ausdrucks *signal-to-noise ratio* – ist das Verhältnis der *Nutzsignalleistung* zur *Rauschleistung*, oder von *signal power* zu *noise power*. Je größer die Nutzsignalleistung im Vergleich zur Rauschleistung, desto leichter kann das Signal entdeckt werden.

$$\text{SVR} = \frac{\text{Nutzsignalleistung}}{\text{Rauschleistung}} = \frac{P_{\text{Signal}}}{P_{\text{Rauschen}}}. \quad (2.41)$$

Dabei ist P_{Signal} die durchschnittliche Leistung des Signals, und P_{Rauschen} die durchschnittliche Leistung des Rauschens. Leistung ist Energie pro Zeiteinheit, $P = E/t$; die Signalenergie ist

$$E_s = \int_{-\infty}^{\infty} |s(t)|^2 dt. \quad (2.42)$$

Nach Parsevals Satz (vergl. Satz A.3.9 im Anhang, Seite 170) gilt dann

$$E_s = \int_{-\infty}^{\infty} |S(\omega)|^2 d\omega, \quad (2.43)$$

wenn $S(\omega)$ die Fouriertransformierte des Signals $s(t)$ ist. Analog dazu ist die Energie des Rauschens durch

$$E_r = \int_{-\infty}^{\infty} |F(\omega)|^2 d\omega \quad (2.44)$$

gegeben, wenn F die Spektraldichtefunktion des Rauschens ist. In Abschnitt A.9 des Anhangs wird gezeigt, dass

$$E_r = \sigma^2 = R(0) \quad (2.45)$$

gerade die durchschnittliche Energie des Rauschens ist, vergl. insbesondere (A.163), Seite 188. Das Signal-Rausch-Verhältnis ist dann durch

$$\text{SVR} = \frac{E_s}{E_r} = \frac{E_s}{\sigma^2} \quad (2.46)$$

definiert, – die Zeiteinheit kürzt sich heraus.

In gängigen Lehrbüchern über Anwendungen der Signalentdeckungstheorie in der Psychophysik wird die in (2.40) eingeführte Größe d' ohne den Umweg über die Definition von $X = \log \lambda(\eta)$ eingeführt; es wird einfach angenommen, dass d' ein Maß für die Sensitivität ist und die Entdeckungswahrscheinlichkeit vom Wert normalverteilter zufälliger Veränderlichen abhängt, die aber nicht notwendig gleiche Varianzen haben müssen (vergl. etwa Thomas (1999, 2003)).

2.3 Spezielle Modelle des Entdeckens

Die Annahme, dass Versuchspersonen sich nach Maßgabe des Likelihood-Quotienten entscheiden, ob ein Stimulus gezeigt wurde oder nicht, ist ein mögliches, aber kein zwingendes Postulat. Insbesondere muß geklärt werden, welche Interpretation dem Term R_T , wie er in (2.32) eingeführt wurde, zugeordnet werden kann. Dass man R_T als Korrelation zwischen dem Signal s und der mittleren Antwort g bezeichnen kann, ist zunächst eine formale Interpretation, denn es bleibt ja offen, wie ein entdeckendes System diese Größe berechnen soll. Eine Möglichkeit besteht darin, s als gespeichertes Bild des Stimulus zu interpretieren. Das unterstellt, dass ein exaktes Bild von s gespeichert werden kann. Eine Möglichkeit, diesen Fall zu interpretieren, besteht in der Annahme, dass ein Filter mit der Impulsantwort

$$h(T - t) = s(t) \quad (2.47)$$

existiert; dann folgt

$$\int_0^T g(t)h(T - t)dt = \int_0^T g(t)s(t)dt. \quad (2.48)$$

Es ist unplausibel, dass für einen beliebigen Stimulus korrespondierende Filter existieren; es ist aber denkbar, dass sich solche Filter im Laufe eines Entdeckungsexperiments bilden. Darauf wird später noch eingegangen.

In der visuellen Psychophysik sind dann auch zwei andere, allerdings miteinander konkurrierende Modelle des Entdeckens populär: *Peak-Detection* und *zeitliche Wahrscheinlichkeitssummation*. Diese beiden Modelle werden im Folgenden vorgestellt.

2.3.1 Peak-Detection

Hier wird angenommen, dass die Wahrscheinlichkeit des Entdeckens durch den Wert des Maximums der deterministischen Antwort bestimmt wird. Der Ausdruck *peak-detection* erklärt sich aus der Definition dieses Kriteriums: wenn der *peak* – also der Maximalwert – einen Wert η_0 erreicht, wird der Stimulus mit der Wahrscheinlichkeit p erreicht. Ein einfaches Modell zur Peak-detection-Annahme besteht in der Annahme der Existenz einer zufälligen Veränderlichen ξ , über die die Wahrscheinlichkeit

$\psi(c)$ des Entdeckens in Abhängigkeit von der Stimulusintensität oder Stimuluskontrast c definiert wird: mit $X = g_0 + \xi$ soll dementsprechend gelten

$$\psi(c) = P(g_0 + \xi > S) = 1 - P(\xi \leq S - g_0(c)), \quad g_0(c) = \max_{t \in [0, T]} g(t; c). \quad (2.49)$$

ψ ist eine Funktion von c , d.h. ψ ist eine *psychometrische Funktion*. ψ wird durch die Wahrscheinlichkeitsverteilung der zufälligen Veränderlichen ξ bestimmt. Formal ist die Definition von ψ einfach, die Frage ist aber, welche Bedeutung ξ hat. So lange man nur an der Charakterisierung der deterministischen Antwort und damit des Systems interessiert, kann man die Frage zunächst unbeantwortet lassen, die Frage, ob das Peak-detection-Kriterium überhaupt gilt, ist zunächst wichtiger. Sie läßt sich empirisch angehen; darauf wird weiter unten noch ausführlich eingegangen.

2.3.2 Wahrscheinlichkeitssummation

Es wird zunächst die zeitliche Wahrscheinlichkeitssummation betrachtet. Der Ausdruck *Wahrscheinlichkeitssummation* reflektiert eigentlich eine begriffliche Ungenauigkeit und ist dementsprechend in der Wahrscheinlichkeitstheorie oder der Statistik nicht üblich, hat sich aber insbesondere in der visuellen Psychophysik eingebürgert und wird nur deswegen hier übernommen. Für den Zeitbereich erklärt er sich durch die Annahme, dass der Stimulus zu irgendeinem Zeitpunkt innerhalb des Versuchsdurchganges entdeckt werden kann. Formal läßt sich das Postulat der Wahrscheinlichkeitssummation formulieren, wenn man die Aktivität als Funktion der Zeit $\eta(t)$ konzipiert, wobei man

$$\eta(t) = g(t) + \xi(t), \quad 0 \leq t \leq T \quad (2.50)$$

schreiben kann. Definiert man die Erwartungsfunktion $E(\eta(t)) = g(t)$, so ist $\xi(t) = \eta(t) - g(t)$ einfach die Abweichung der Aktivität von der mittleren Aktivität, jeweils zur Zeit t . Sicherlich erreicht $\eta(t)$ mindestens einmal während der Versuchsdurchganges, repräsentiert durch das Intervall $[0, T]$, den Wert η_0 , wenn $\max_{t \in [0, T]} \eta(t) > S$. Dann ist die psychometrische Funktion durch

$$\psi(c) = P(\max_{t \in [0, T]} \eta(t) > S | c) = 1 - P(\max_{t \in [0, T]} \eta(t) \leq S | c) \quad (2.51)$$

definiert. Das Problem, einen Ausdruck für die Wahrscheinlichkeit $P(\max_{t \in [0, T]} \eta(t) \leq S)$ zu finden, ist schwierig. Es läßt sich aber eine allgemeine Form für ψ anschreiben, die sich aus der Tatsache ergibt, dass die Frage nach der Verteilung von $\max_{t \in [0, T]} \eta(t)$ als Frage nach der Verteilung einer Wartezeit τ formuliert: τ ist die Zeit, die verstreicht, bis α zum ersten Mal den Wert η_0 erreicht. Für $\tau \leq T$ wird der Stimulus entdeckt, für $\tau > T$ nicht, d.h.

$$\psi(c) = P(\max_{t \in [0, T]} \eta(t) > S | c) = P(\tau \leq T | c). \quad (2.52)$$

Man kann nun die bedingte Wahrscheinlichkeit

$$P(t, t + \Delta t) = \frac{P(\tau \in [t, t + \Delta t])}{P(\tau > t)} \quad (2.53)$$

betrachten; dies ist die Wahrscheinlichkeit, dass α im Intervall $[t, t + dt)$ den Wert η_0 erreicht, unter der Bedingung, dass der Wert η_0 im Intervall $[0, t)$ nicht erreicht wurde. Für beliebig kleines $\Delta t = dt$ wird aus $P(t, t + \Delta t)$

$$P(t, t + dt) = \frac{f(t)dt}{1 - F(t)} = \phi(t)dt; \quad (2.54)$$

$\phi(t) = f(t)/(1 - F(t))$ heißt *Hazardfunktion*, und F ist die Verteilungsfunktion der Wartezeit, f die zugehörige Dichtefunktion $f(t) = dF(t)/dt$. Es läßt sich zeigen, dass ψ dann die Form

$$\psi(c) = P(\tau \leq T) = 1 - \exp \left[- \int_0^T \phi(t; c) dt \right] \quad (2.55)$$

annimmt. Die psychometrische Funktion hat im Falle von Entdecken durch Wahrscheinlichkeitssummation stets diese Form, wird also stets durch ein Integral über das Intervall, das einen Versuchsdurchgang repräsentiert, bestimmt. Das Problem besteht nun darin, die Hazardfunktion ϕ zu bestimmen.

Die Annahme der Wahrscheinlichkeitssummation erscheint zunächst als plausibel: es soll ja nur darauf ankommen, dass die Aktivität zu irgendeinem Zeitpunkt während des Versuchsdurchganges einen kritischen Wert erreicht, damit der Stimulus entdeckt wird. Die Peak-Detection-Annahme erscheint dagegen als artifiziell, – denn es ist nicht klar, warum es nur auf den Wert von $g_0 = \max_t g(t)$ ankommen soll. Diese Annahme hat aber den Vorteil, dass sie eine ziemlich direkte Erfassung der deterministischen Antwort g des entdeckenden Systems erlaubt. Der Ausdruck (2.89) erlaubt es überdies, empirisch zu testen, ob Entdecken durch Peak-Detection oder durch Wahrscheinlichkeitssummation geschieht; darauf wird weiter unten noch eingegangen. Zunächst wird die Systemidentifikation über die Annahme von Peak-Detection illustriert.

Das Modell der Wahrscheinlichkeitssummation kann auch auf den Ortsbereich übertragen werden, bzw. auf Modelle, denen zufolge ein Muster durch Aktivierung verschiedener "Kanäle" – etwa Ortsfrequenzfilter – repräsentiert wird. Bei diesen Modellen wird dann im Allgemeinen stillschweigend zeitliche Peak-Detektion vorausgesetzt und als Annahme auch nicht weiter überprüft. Der Grund hierfür ist wohl, dass dadurch das jeweils betrachtete Entdeckensmodell sehr vereinfacht wird, – eine zusätzliche Berücksichtigung von zeitlicher Wahrscheinlichkeitssummation würde die Komplexität des Modells sehr erhöhen, zumal nach nicht ausgeschlossen werden kann, dass verschiedene Ortskanäle verschiedene zeitliche Charakteristika haben. Kanäle für niedrige Ortsfrequenzen könnten sehr viel schneller reagieren als Kanäle für höhere Ortsfrequenzen.

2.4 Psychometrische Funktionen

2.4.1 Allgemeine Modelle

Ein großer Teil der Versuche, sensorische Systeme systemtheoretisch zu charakterisieren, bezieht sich aus Daten aus Entdeckensexperimenten: Stimuli werden mit scha-

cher Intensität bzw. schwachem Kontrast c gezeigt und die Aufgabe der Versuchsperson ist, den Stimulus zu entdecken (Detektion), oder es werden sich nur wenig unterscheidende Stimuli präsentiert und die Aufgabe besteht darin, den Unterschied zu entdecken (Diskrimination). Die Stimuli werden dann variiert und die Intensität bzw. der Kontrast werden für jede Variation so bestimmt, dass die Wahrscheinlichkeit $\psi(c)$ des Entdeckens des Stimulus oder des Unterschieds zwischen Stimuli konstant bleibt.

Definition 2.4.1 *Es sei $\psi(c)$ die Wahrscheinlichkeit des Entdeckens eines Stimulismusters, oder $\psi(\Delta c)$ die Wahrscheinlichkeit des Entdeckens eines Unterschieds zwischen Stimulismustern, wobei c der Kontrast des Musters bzw. Δc der Unterschied der Kontraste der Stimulismuster ist, und $c \in \mathbb{R}$ bzw. $\Delta c \in \mathbb{R}$. Die Abbildung $\psi : \mathbb{R} \rightarrow [0, 1]$ heißt psychometrische Funktion.*

Dementsprechend gilt für $c \geq 0$ stets $0 \leq \psi(c) \leq 1$; ebenso gilt $0 \leq \psi(\Delta c) \leq 1$, wobei aber nicht $\Delta c \geq 0$ gelten muß. Für die im Folgenden behandelten Fälle gilt, dass ψ monoton mit c wächst, also $\psi(c) \leq \psi(c')$, wenn $c < c'$. Weiterhin findet man i. A., dass ψ sigmoidal, also irgendwie S-förmig ist. Gilt $\psi(c_p) = p$, so heißt c_p die Schwellenintensität oder der Schwellenkontrast bezüglich p . Üblicherweise wird $p = .5$ oder $p = .75$ gewählt. Es bedeutet keine Einschränkung der Allgemeinheit, ψ als stetige Funktion von c zu konzipieren, d.h. den Fall

$$\psi(c) = \begin{cases} 0, & c < c_0 \\ 1, & c \geq c_0. \end{cases} \quad (2.56)$$

auszuschließen; empirisch ist dieser Fall wohl kaum jemals aufgetreten. Demnach sind die Reaktionen des Entdeckens und Nichtentdeckens stochastisch mit den Werten von c verknüpft. Die Frage ist, wie diese Verknüpfung zu konzipieren ist.

Die üblicherweise gefundene S-Förmigkeit legt zunächst nahe, die Gauß-Verteilung an die Daten anzupassen, d.h. die psychometrische Funktion über die Gauß-Verteilung zu definieren. Dementsprechend werden die stochastischen Effekte als normalverteilt angenommen. Damit wird postuliert, dass die durch den Stimulus erzeugte Aktivierung durch eine zufällige Veränderliche X repräsentiert wird mit einem Erwartungswert $\mu = \mathbb{E}(X)$, der wiederum eine Funktion von c ist, also $\mu = \mu(c)$. Die nächst einfache Annahme ist dann, μ also proportional zu c anzusetzen, also $\mu = c\mu_0$, wobei nun μ_0 eine Konstante ist, die von der experimentellen Situation abhängt, also von den Bedingungen, unter denen der Stimulus präsentiert wird (Adaptationszustand der Versuchsperson, aufmerksamkeitssteuernde Instruktionen, Elemente wie flankierende Linien, die wahrnehmungsmäßig mit dem zu entdeckenden Stimulus in Wechselwirkung stehen, etc.). Die psychometrische Funktion hat dann die Form

$$\psi(c) = 1 - \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^S \exp \left[-\frac{(x - \mu(c))^2}{2\sigma^2} \right] dx = 1 - \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{z_S\mu(c) + \sigma} e^{-z^2/2} dz, \quad (2.57)$$

oder kurz

$$\psi(c) = 1 - \Phi(z_S\mu(c) + \sigma), \quad z_S = (S - \mu(c))/\sigma, \quad (2.58)$$

wobei Φ wie üblich die Verteilungsfunktion einer $N(0, 1)$ -verteilten zufälligen Veränderlichen ist. Dieser Ausdruck kann noch verallgemeinert werden, indem man die Möglichkeit, dass die Versuchsperson korrekt rät, dass ein Reiz präsentiert wurde, obwohl sie ihn nicht wahrgenommen hat. Ist γ die Wahrscheinlichkeit, korrekt zu raten, kommt es mit einer Wahrscheinlichkeit $1 - \Phi$ zu einer korrekten Antwort, dass ein Stimulus gezeigt wurde, wenn die Versuchsperson den Stimulus tatsächlich wahrgenommen hat, oder mit der Wahrscheinlichkeit Φ hat sie ihn nicht wahrgenommen, rät aber korrekt mit der Wahrscheinlichkeit γ , so dass

$$P(\text{"ja"}) = 1 - \Phi + \gamma\Phi = 1 - (1 - \gamma)\Phi$$

gilt, d.h.

$$\psi(c) = 1 - (1 - \gamma)\Phi(z_S\mu(c) + \sigma), \quad z_S = (S - \mu(c))/\sigma. \quad (2.59)$$

Man kann den Faktor $1 - \gamma$ als Ratekorrektur bezeichnen; γ ist ein freier Parameter, der zusätzlich aus den Daten geschätzt werden muß. Nimmt man $\mu = c\mu_0$ an, so sind die übrigen freien Parameter σ und μ_0 . Unabhängig von der Methode der Bestimmung der psychometrischen Funktion (Methode der konstanten Stimuli, Grenzmethode (method of limits), up-down-Methode, etc) wird man auf die Probit-Analyse zur Schätzung der Parameter geführt.

Die Frage nach einer Begründung für die Wahl einer normalverteilten zufälligen Veränderlichen X als Repräsentation der Aktivierung durch den Stimulus wird üblicherweise durch den globalen Hinweis auf den Zentralen Grenzwertsatz beantwortet, demzufolge die Summe vieler, voneinander unabhängiger Effekte zu einer normalverteilten Größe führt. Da die Daten häufig mit dem Ansatz (2.59) kompatibel sind und man nicht primär an den stochastischen Effekten, sondern an der Funktion $\mu(c)$ interessiert ist, da sie mit den Eigenschaften des System in Beziehung gesetzt wird, wird dieser Hinweis als hinreichende Motivation für die Definition (2.59) von ψ akzeptiert. In Präcomputerzeiten kann aber die Berechnung von Φ eine lästige Aufgabe sein, die eine Suche nach Alternativen erzeugt. Eine rechnerisch angenehme Form, die von der Gauß-Verteilung praktisch nicht zu unterscheiden ist, ist die logistische Verteilung, die auf die Definition

$$\psi(c) = \frac{1}{1 + \exp\left(-\frac{(X - \mu(c))\sqrt{\pi}}{\sigma}\right)} = \frac{1}{1 + \exp(\alpha\mu(c) + \beta)} \quad (2.60)$$

führt, wobei der rechte Ausdruck einfach eine Reparametrisierung der ursprünglichen Form ist. Die Ratekorrektur ist hier der Einfachheit halber fortgelassen worden.

Im Zusammenhang mit der Hypothese, dass visuelle Muster durch Zusammenschaltung neuronale Kanäle, die für verschiedene Ortsfrequenzen spezifisch sind, wahrgenommen werden, wurde die These aufgestellt, dass diese Ortsfrequenzkanäle unabhängig voneinander arbeiten; im Rahmen eines Entdeckensexperiments sollte dies insbesondere stochastische Unabhängigkeit bedeuten. Sachs, Nachmias und Robson (1971) formulierten dementsprechend die Hypothese, dass ein Stimulismuster dann entdeckt wird, wenn mindestens einer von insgesamt n solcher Kanäle, die durch das Muster aktiviert werden, eine hinreichend große Aktivität zeigt. Dementsprechend soll dann für jeden Kanal eine zufällige Veränderliche X_j , $j = 1, \dots, n$

definierbar sein, und die psychometrische Funktion ist nun definiert durch

$$\psi(c) = P(\max(X_1, \dots, X_n) > S), \quad (2.61)$$

wobei S ein bestimmter, minimaler Wert ist, der erreicht werden muß, damit die Aktivierung als indikativ für die Präsentation eines Stimulus interpretiert wird. Nun ist sicherlich

$$\psi(c) = 1 - P(\max(X_1, \dots, X_n) \leq S),$$

und die Annahme der stochastischen Unabhängigkeit bedeutet

$$\{\max(X_1, \dots, X_n) \leq S\} = \{X_1 \leq S \cap X_2 \leq S \cap \dots \cap X_n \leq S\} = \prod_{j=1}^n P(X_j \leq S),$$

so dass die psychometrische Funktion durch

$$\psi(c) = 1 - \prod_{j=1}^n P(X_j \leq S) \quad (2.62)$$

definiert ist. $P(X_j \leq S)$ ist eine Verteilungsfunktion, die für den j -ten Kanal charakteristisch ist. Die einfachste Annahme ist nun, dass die Verteilungen der X_j alle vom gleichen Typ sind und sich allenfalls durch die Werte der Parameter dieser Verteilungen unterscheiden. Sachs et al. (1971) gingen von der üblichen Annahme der Gauß-Verteilung der X_j aus, so dass die psychometrische Funktion insbesondere die Form

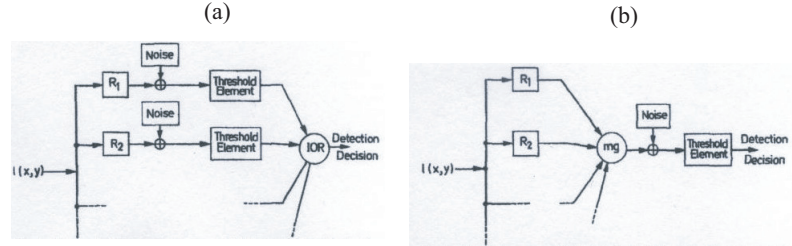
$$\psi(c) = 1 - \prod_{j=1}^n \Phi(z_S \mu_j(c) + \sigma_j) \quad (2.63)$$

hat. Weitere Vereinfachungen ergeben sich, wenn man insbesondere $\mu_j(c) = c\mu_{0j}$ und $\sigma_j = \sigma$ für alle j postuliert und $\mu_{0j} = |H_j(\omega_j)|$ setzt. Während die Daten von Sachs et al. insgesamt mit der Hypothese der W-Summation zwischen Kanälen kompatibel erschienen, ergaben sich eine Reihe von neuen Fragen, – etwa nach der Bandbreite der Ortsfrequenzkanäle, und wie die Hypothese der stochastischen Unabhängigkeit, die ja eigentlich nur der Vereinfachung dienen soll, damit verträglich ist, dass das Ortsfrequenzspektrum eines Musters ja im Allgemeinen eine stetige Funktion der Ortsfrequenzen $\omega = 2\pi f$ ist und der j -te Kanal wohl auch für der Ortsfrequenz ω_j benachbarte Frequenzen empfindlich ist, da sonst angenommen werden müßte, dass die Systemfunktionen der einzelnen Kanäle durch Dirac-Delta-Funktionen definiert sind, was extrem unplausibel ist. Rechnerisch erweisen sich sowohl die Gauß- als auch die logistische Verteilung als umständlich, und so liegt es nahe, nach rechnerisch einfacheren Alternativen zu suchen.

2.4.2 Das Quick-Modell: W-Summation zwischen Kanälen

In einem kurzen Artikel hat Quick (1974) diese Alternative geliefert. Quick verweist zunächst auf eine Arbeit von Abadi und Kulikowski (1973), denen zufolge die (hypothetischen) Ortsfrequenzkanäle vor dem neuronalen Ort der Entscheidung einerseits

Abbildung 2.2: Quicks Schemata: (a) W-Summation, (b) Nichtlineare Summation (Vector-magnitude-Modell. Aus Quick (1974).



in guter Näherung linear und unabhängig voneinander sind. Die 'Nettoaktivierung', wie Quick die Aktivierung nennt, bezeichnet Quick mit R_j , um dann eine *magnitude function* zu definieren:¹

$$\|\vec{R}\| = \left[\sum_j (R_j(s))^\beta \right]^{1/\beta}. \quad (2.64)$$

$R_j(s)$ ist die Antwort des j -ten Kanals auf den Stimulus s . Quick bezeichnet das n -Tupel $(R_1, \dots, R_n)$ als Vektor $\vec{R}$; für $\beta = 2$ ist $\|\vec{R}\|$ gerade die übliche Länge dieses Vektors. Für $\beta \neq 2$ kann man $\|\vec{R}\|$ die verallgemeinerte Vektorlänge nennen. Dies motiviert den Ausdruck *vector-magnitude model*, den Quick für sein Modell gewählt hat.

Für beliebiges $\beta \in \mathbb{R}$ repräsentiert einen Ausdruck, der formal ähnlich einem Ausdruck ist, den Minkowski im Zusammenhang mit Fragen der Relativitätstheorie (Raumkrümmung) diskutiert hat und der in der Multidimensionalen Skalierung als Minkowski-Metrik auftritt, wobei die R_j durch Distanzen d_j auf der j -ten latenten Dimension, hinsichtlich der Stimuli beurteilt werden, ersetzt werden. Deswegen wird auch gelegentlich in Bezug auf (2.64) von der Minkowski-Summation geredet.

Die Frage ist nun, wie $\|\vec{R}\|$ zur Wahrscheinlichkeit des Entdeckens in Beziehung gesetzt wird. Wenn nun die Verteilung des Rauschens derart sei, dass

$$\psi_Q(c) = 1 - 2^{-\|R\|^\beta} \quad (2.65)$$

gelte, so Quick, so könne man auch

$$\psi_Q(c) = 1 - \prod_{j=1}^n [1 - (1 - 2^{-R_j^\beta})] \quad (2.66)$$

schreiben (der Index Q soll das Quick-Modell indizieren). Damit ist das Vektorgrößenmodell (2.65) auch ein W-Summationsmodell. Abbildung 2.2 zeigt die Sche-

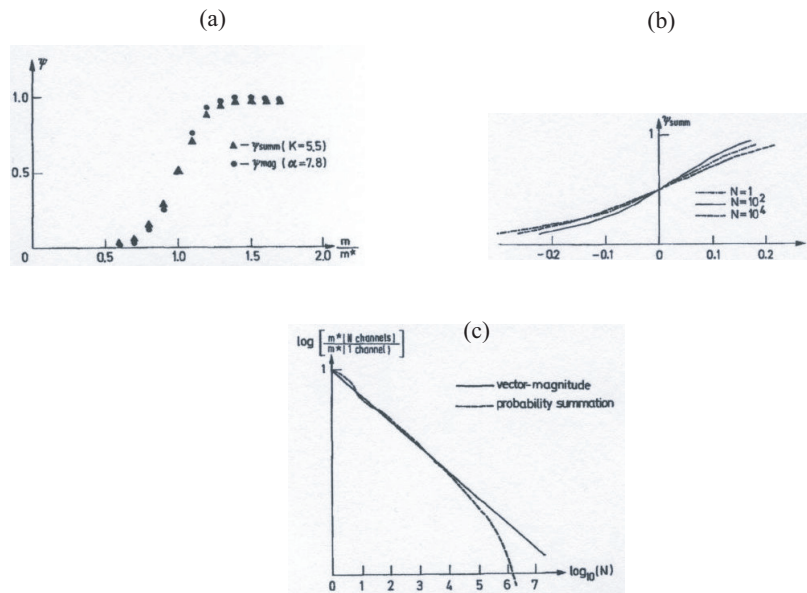
¹Quick schreibt α statt β . In den Anwendungen des Quick-Modells hat sich dann die hier gewählte Schreibweise – also β für den Exponenten – durchgesetzt. Um die Diskussion der Bedeutung des Parameters in der Literatur nicht zu verwischen, ist hier die jetzt gebräuchliche Schreibweise benutzt worden.

mata für W-Summation (a), und das Magnitude-Modell (b). Eine etwas direktere Weise, ψ_Q anzuschreiben, ist natürlich die zu (2.65) äquivalente Form

$$\psi_Q(c) = 1 - \prod_{j=1}^n 2^{-R_j^\beta}. \quad (2.67)$$

Die Modelle unterscheiden sich in Bezug auf den Ort des neuronalen Rauschens: im W-Summationsmodell hat jeder Kanal eine eigene Rauschquelle, und der Entdeckensmechanismus ist eine ODER-Schaltung, während im Summations- oder Magnituden-Modell die Aktivitäten der individuellen Kanäle zuerst deterministisch in einer UND-Schaltung zusammengefaßt werden, und erst die zusammengefasste Aktivität wird durch Rauschen gestört. Abb. 2.3 zeigt einen Vergleich des W-Summationsmodells

Abbildung 2.3: (a) Vergleich der Vorhersagen des W-Summationsmodell von Sachs et al. (1971) und dem Vektor-Magnituden-Modell von Quick (1974), (b) Verlauf von ψ_Q für verschiedene Werte von N . Nach Quick (1974).



von Sachs et al. (1971) mit dem Quickschen Vektor-Magnituden-Modell. Obwohl beim Modell von Sachs et al. Gauß-Verteilungen angenommen werden und das Magnituden-Modell durch die zunächst nicht weiter einsichtigen Annahme (2.65) charakterisiert wird, ist die Übereinstimmung bemerkenswert gut. Erst bei einer Anzahl von 10^4 Kanälen zeigen sich Abweichungen der Vorhersagen der beiden Modelle (vergl. Abb. 2.3), (c).

Quick gibt nicht an, welche Betrachtungen ihn auf sein Modell geführt haben, und insbesondere die Wahl der Funktion $2^{-\|\vec{R}\|}$ in (2.65) bleibt unerläutert, und dementsprechend werden auch keine theoretischen Gründe für die Übereinstimmung

des Sachs et al.schen und des Quickschen Modells, wie sie in Abb. 2.3), (a) zum Ausdruck kommt, diskutiert. Eine Auflösung des Rätsels ergibt sich aber, wenn man den Term $2^{-\|R\|^\beta}$ in (2.65) betrachtet. Nimmt man den natürlichen Logarithmus dieses Terms, so erhält man

$$\log 2^{-\|R\|^\beta} = \|R\|^\beta \log 2,$$

und dementsprechend

$$2^{-\|R\|^\beta} = \exp\left(-\nu\|R\|^\beta\right), \quad \nu = \log 2.$$

Ersetzt man also in (2.65) den Ausdruck $2^{-\|R\|^\beta}$ durch $\exp(-\nu\|R\|^\beta)$, so erhält man

$$\psi_Q(c) = 1 - \exp\left(-\nu\|R\|^\beta\right). \quad (2.68)$$

Quick hat also implizit die Weibull-Funktion gewählt, wie es scheint, ohne sich dessen bewußt gewesen zu sein. Macht man noch von der Linearitätsannahme Gebrauch, so läßt sich

$$R_j(c) = c R_{0j}$$

schreiben, wobei R_{0j} die Antwort des j -ten Kanals auf die Intensität 1 ist. Dann läßt sich (2.68) in der Form

$$\psi_Q(c) = 1 - \exp\left(-\nu\|c R_0\|^\beta\right)$$

schreiben, und absorbiert man ν in den Term $\|R_0\|$, so erhält man die psychometrische Funktion

$$\psi_Q(c) = 1 - \exp\left[-\|R_0\|^\beta c^\beta\right], \quad \|R_0\| = \left[\sum_j R_{0j}^\beta\right]^{1/\beta} \quad (2.69)$$

oder

$$\psi_Q(c) = 1 - (1 - \gamma)e^{-\alpha c^\beta}, \quad \alpha = \|R_0\|^\beta \quad (2.70)$$

wenn man die Möglichkeit des korrekten Ratens berücksichtigen will. Der Einfachheit halber wird im Folgenden Bezug auf (2.69) genommen.

Anmerkung: Quick hat bei seinen Simulationen die Maxima von normalverteilten Variablen betrachtet. Die Grenzverteilung des Maximums normalverteilter Variablen ist aber gar nicht die Weibull-Verteilung, sondern die *Extremwertverteilung* oder auch *Doppelte Exponentialverteilung*

$$G(X \leq x) = \exp(-e^{-x}). \quad (2.71)$$

Dass die Weibull-Verteilung so gut für die Maxima von Gauß-verteilten Variablen passt, wie in Abb. 2.70 (c) angedeutet, zeigt, wie gut die Weibull-Verteilung an Daten angepasst werden kann. Weitere Beispiele findet man in Mortensen (2002). Auch die in Gleichung (2.89), Seite

107 gegebene psychometrische Funktion, der die Annahme eines Gauß-Prozesses mit nachgerade beliebiger Autokorrelation zugrundeliegt, kann durch eine Weibull-Verteilung approximiert werden. Inspektion von (2.89) zeigt übrigens, dass sie mit der doppelten Exponentialverteilung (2.71) verwandt ist. Diese Verwandtschaft resultiert eben aus der Tatsache, die die Maxima korrelierender Gauß-Variablen betrachtet werden.

Das Quick-Modell ist in der Form (2.69) zum Standardmodell in der visuellen Psychophysik geworden, gelegentlich in seiner Form als Weibull-Funktion. Die Hauptgründe dafür scheinen (i) seine Einfachheit und (ii) der im Allgemeinen gute Fit für die verschiedensten Daten aus Entdeckensexperimenten zu sein. Darüber hinaus wird darauf hingewiesen, dass die Weibull-Funktion in der Reliabilitätstheorie, in der die Zuverlässigkeit von Materialien bzw. von aus vielen Komponenten bestehenden Apparaten in Abhängigkeit von der Zuverlässigkeit der Komponenten untersucht wird, eine zentrale Rolle spielt, wobei angemerkt wird, dass insbesondere die Wahrscheinlichkeit des Brechens von Ankerketten – die ja brechen, wenn das schwächste Glied bricht – durch die Weibull-Funktion gegeben ist. Wie das Paradigma brechender Ankerketten auf das Entdecken von Reizen zu übertragen ist, bleibt eine offene Frage, deren Beantwortung der Kreativität des Lesers überlassen bleibt. Man mag das Paradigma in die Sprache des Entdeckens durch unabhängige neuronale Kanäle übersetzen: ein Muster wird entdeckt, wenn der Damm, der die Aktivität der Kanäle begrenzt, bei mindestens einem Kanal bricht. Ob diese Analogie hilfreich ist, ist eine andere Frage.

Die psychometrische Funktion (2.69) definiert ein Hochschwellenmodell. Für $c = 0$ ist $\|R_0\|c^\beta = 0$ und $\psi_Q(c) = 0$; falsche Alarme können nur erklärt werden, wenn man die Ratewahrscheinlichkeit γ wie in (2.70) einführt. Nachmias (1981) diskutiert die Schätzungen der freien Parameter und deren Deutung, ohne auf die Möglichkeit, das Hochschwellenmodell, aufzugeben, hinzuweisen. Die Problematik der Annahme stochastisch unabhängiger Kanäle oder benachbarter retinaler Positionen (*spatial probability summation*) wird dabei ebensowenig wie die Hochschwellenannahme sowie der ebenfalls implizit geforderte Rektifikator – die Antworten der Kanäle gehen ja betragsmäßig, also in der Form $|g_j|$, in die psychometrische Form ein – nie thematisiert, – man denke an die Modelle von Graham und Rogowitz (1976), Graham (1977), Wilson und Giese (1977), Wilson (1978), Wilson und Bergen (1979), Wilson, Philipps, Rentschler und Hilz (1979), Wilson (1980), um nur einige der "früheren" Arbeiten zu nennen, in denen nicht nur Ortsfrequenzkanäle, sondern auch *spatial probability summation*, also W-Summation zwischen benachbarten retinalen Positionen anhand des Quickschen Ansatzes modelliert wird. xxx

Zumindest das Hochschwellenmodell sowie die dadurch erzwungene Forderung nach einem Rektifikator sind aber nicht notwendig, wenn man von der Weibull-Verteilung ausgeht. Die Weibull-Verteilung wird in Anhang A.10, Seite 194, behandelt. Geht man, will man W-Summation zwischen stochastisch unabhängigen Kanälen modellieren, davon aus, dass ein Kanal K_j die Präsentation eines Stimulus signalisiert, wenn die durch die zufällige Veränderliche X_j Weibull-verteilt ist und einen Wert größer als S annimmt, wenn entdeckt wird, ist der Typ (A.181), Seite

194, indiziert:

$$P(X_j \leq S) = \exp \left[- \left(\frac{\eta_j - S}{\sigma_j} \right)^\beta \right], \quad X_j \leq \eta_j. \quad (2.72)$$

Es sei noch einmal darauf hingewiesen, dass die X_j nun als auf dem Intervall $(-\infty, \eta_j]$ verteilt angenommen werden; η_j kann vielleicht als maximale Spike-Rate interpretiert werden. Mit $X_j = g_j(c) + \xi_j$ erhält man dann

$$P(X_j \leq S) = P(\xi_j \leq S - g_j(c)) = \exp \left[- \left(\frac{\eta_j - (S - g_j(c))}{\sigma_j} \right)^\beta \right], \quad S - g_j(c) \leq \eta_j$$

und daraus die psychometrische Funktion

$$\psi(c) = 1 - \exp \left[- \sum_j \left(\frac{\eta_j - (S - g_j(c))}{\sigma_j} \right)^\beta \right]. \quad (2.73)$$

Wenn es Sinn macht, $\eta_j = \eta$ und $\sigma_j = \sigma$ für alle j anzunehmen, kann $1/\sigma$ in η , S und $g_j(c)$ absorbiert werden; insbesondere mit $g_j(c) = cg_{0j}$ erhält man dann

$$\psi(c) = 1 - \exp \left[- \sum_j (\eta^* - cg_{0j})^\beta \right], \quad \eta^* = \eta - S. \quad (2.74)$$

Für $\eta^* > 0$ sind echte falsche Alarme möglich und die Annahme eines Rektifikators wird unnötig.

Die Problematik eines Hochschwellenmodells kann also vermieden werden, wenn man Weibull-verteilte zufällige Variable zur Repräsentation der Aktivitäten der Kanäle annehmen möchte. Es bleibt die Frage, warum man überhaupt die Weibull-Verteilung für die Aktivitäten annehmen möchte. Was benötigt wird ist ein Modell – also Grundannahmen – der Aktivierung, dass die Weibull-Verteilung impliziert. Es ist gegenwärtig kein solches überprüfbares Modell bekannt: Modelle der Aktivierung von Neuronenpopulationen können häufig als Diffusionsmodelle formuliert werden und implizieren dann die Gauß-Verteilung. Man kann natürlich einen rein pragmatischen Standpunkt einnehmen und darauf verweisen, dass Modelle, die auf der (ad-hoc-)Annahme der Weibull-Verteilung basieren, im Allgemeinen gut an die Daten anzupassen sind, was für viele Fragestellungen auch genügt, sofern die Interpretation der Daten nicht von der postulierten Verteilungsfunktion der zufälligen Veränderlichen abhängt. Vielfach werden allerdings Aussagen über das sensorische oder neuronale Rauschen aus dem Wert des Exponenten β abgeleitet, dessen Wert eben als indikativ für das Ausmaß des Rauschens gilt: ein großer Wert von β impliziert eine steile psychometrische Funktion, und da darüber hinaus die Varianz der zufälligen Veränderlichen umgekehrt proportional zum Wert von β ist, wird gefolgert, dass ein großer Wert von β ein geringes Rauschen bedeutet (so etwa Watson (1981), Blommaert und Roufs (1987)). Ob dieser Schluß gerechtfertigt ist, bleibt allerdings zu diskutieren, ebenso wie die Frage, in welcher Weise eine zufällige Veränderliche denn überhaupt die zeitlich ausgedehnte Aktivierung eines neuronalen Kanals repräsentieren kann

Problematisch ist die Annahme der stochastischen Unabhängigkeit zwischen Kanälen oder benachbarten spatialen Positionen. Zur Illustration wird zunächst die Anwendung des Quickschen Modells auf die zeitliche W-Summation betrachtet.

2.4.3 Zeitliche Wahrscheinlichkeitssummation

Dass zeitliche W-Summation im Entdeckensprozess eine Rolle spielen könnte, kann anhand vieler Daten vermutet werden. Nimmt man zeitliche W-Summation an, so erhält die Repräsentation der Aktivität durch eine zufällige Veränderliche sofort Sinn. Denn es wird ja nun gefordert, dass der Stimulus entdeckt wird, wenn

$$X_j^+ = \max_{t \in [0, T]} X_j(t) > S,$$

damit der j -te Kanal den Stimulus entdeckt; dabei ist $X_j(t)$ eine Funktion der Zeit, insbesondere ist $X_j(t)$ die Trajektorie eines zufälligen Prozesses. Die zufällige Veränderliche, die die Aktivität des j -ten Kanals repräsentiert, ist jetzt X_j^+ . Dabei kann man

$$X_j(t) = g_j(t) + \xi_j(t), \quad t \in [0, T] \quad (2.75)$$

setzen, wobei $[0, T]$ ein Zeitintervall ist, dass einen Versuchsdurchgang der Dauer T repräsentiert, g_j ist die mittlere Aktivierung, wie sie durch den Stimulus erzeugt wird, und $\xi_j(t)$ ist eine Trajektorie des Rauschprozesses im j -ten Kanal.

Die Aufgabe ist nun, die Verteilungsfunktion der X_j^+ bzw. der $\xi_j(t)$ zu bestimmen. Ein Ansatz, diese Aufgabe zu lösen, besteht im ersten Schritt darin, die Aufgabe ein wenig umzuformulieren. Der Stimulus wird, so die Annahme der zeitlichen W-Summation, durch den j -ten Kanal entdeckt, wenn

$$X_j(t) = g_j(t) + \xi_j(t) > S$$

für mindestens ein $t \in [0, T]$, d.h. wenn

$$\xi_j(t) > S - g_j(t) \text{ für mindestens ein } t \in [0, T]. \quad (2.76)$$

Der index j wird im Folgenden unterdrückt, um die Notation zu vereinfachen, die Betrachtung gilt ja für alle Kanäle. Man teilt nun das Intervall $[0, T]$ in n Teilintervalle I_k der Länge T/n , $k = 1, \dots, n$. Man definiert dann eine zufällige Veränderliche, die das Maximum von $\xi(t)$ im k -ten Teilintervall I_k repräsentiert:

$$\xi_k^+ = \max_{t \in I_k} \xi(t), \quad (2.77)$$

bestimmt dann die Verteilung des Maximums der $\xi_1^+, \dots, \xi_n^+$, und läßt anschließend $n \rightarrow \infty$ und damit die Länge der I_k gegen Null gehen. Man wird auf diese Weise auf eine Aufgabe der Extremwertstatistik geführt, für die sich relativ einfache Lösungen finden lassen, sofern man annehmen kann, dass die ξ_k^+ auch für $n \rightarrow \infty$ stochastisch unabhängig sind. Diese Annahme ist allerdings problematisch, denn sie bedeutet, dass $\xi(t)$ eine Trajektorie eines Prozesses ist, der als Weißes Rauschen bekannt ist.

Für Prozesse dieser Art ist die Autokorrelation durch eine Dirac-Delta-Funktion gegeben:

$$R(\tau) = \sigma^2 \delta(\tau). \quad (2.78)$$

Die Fourier-Transformierte von R ist bekanntlich eine Konstante, d.h.

$$F[R(\tau)] = H(\omega) = K \text{ eine Konstante für alle } \omega \quad (2.79)$$

was den Ausdruck *Weißes Rauschen* erklärt², und da die Energie des Prozesses durch

$$E = \int_{-\infty}^{\infty} |H(\omega)|^2 d\omega = \int_{-\infty}^{\infty} K^2 d\omega = \infty \quad (2.80)$$

hat man einen Prozess mit unendlicher Energie, was physikalisch unmöglich ist. Die Unmöglichkeit resultiert natürlich aus der Forderung (2.78), die ja beliebig schnelle Änderungen der Funktion $\xi(t)$ bedeutet, was nicht möglich ist³. Nun führt die Annahme des Weißes Rauschens in der Signaltheorie häufig zu keinen Problemen, und so kann gefragt werden, warum sie im gegenwärtigen Zusammenhang problematisch sein soll. Der Punkt ist, dass Weißes Rauschen in den genannten Anwendungen als Approximation verwendet wird, wenn von vornherein klar ist, dass der Frequenzbereich beschränkt ist, wenn also $|H(\omega)| \approx K$ einerseits und $|H(\omega)| = 0$ für alle $\omega > \omega_0 < \infty$ gilt. Diese Einschränkung wird aber nicht gemacht, wenn die Forderung (2.78) aufgestellt wird, damit man von der Unabhängigkeit der ξ_k^+ für $n \rightarrow \infty$ ausgehen kann, denn nur dann kann man ohne weitere Zusatzbetrachtungen die Grenzwertsätze der Extremwertstatistik Gebrauch machen. Dass man in der Tat zu absurden Folgerungen gelangt, wird weiter unten noch verdeutlicht. Vorher soll die Lösung von Watson (1979) genannt werden, da sie zur Interpretation von experimentellen Ergebnissen, in denen der Effekt der zeitlichen W-Summation eine Rolle zu spielen scheint, ungeprüft zur Interpretation der Daten herangezogen wird.

Watsons Modell Watson (1979) ging ebenfalls von der Aufteilung (2.77) aus, verwies auf das Quick-Modell und nahm deshalb an, dass die Wahrscheinlichkeit, dass der Stimulus im k -ten Intervall *nicht* entdeckt wird, durch

$$1 - p_k = \exp\left(-|g(t_k)|^\beta\right) \quad (2.81)$$

gegeben ist, wobei er sich auf die Version (2.69) des Quickschen Modells bezog. Dann ist die Wahrscheinlichkeit, dass in mindestens einem der n Intervalle der Stimulus entdeckt wird, durch

$$\psi_W(c) = 1 - (1 - \gamma) \prod_{k=1}^n \exp\left(-|g(t_k)|^\beta\right) = 1 - \exp\left[-\sum_k |g(t_k)|^\beta\right] \quad (2.82)$$

gegeben, wenn man annimmt, dass die Entdeckensereignisse in den einzelnen Intervallen I_k *alle* – also auch die unmittelbar benachbarten – stochastisch unabhängig

²Im weißen Licht sind bekanntlich alle Frequenzen bzw. Wellenlängen vertreten.

³Wie man seit Einstein weiß.

sind. Die Ratewahrscheinlichkeit wird hier berücksichtigt, um die Übereinstimmung mit den Watsonschen Formeln aufrechtzuerhalten. Die Summe über endlich viele Intervalle bereitet computationale Mühsal, die Watson überwindet, indem er (i) $n \rightarrow \infty$ und damit $|I_k| \rightarrow 0$ für alle k fordert ($|I_k|$ ist die Länge des k -ten Intervalls). Watson äußert die Ansicht, dass man einfach zum Integral übergehen könne und offeriert deshalb die Formel

$$\psi_W(c) = 1 - (1 - \gamma) \exp \left[- \int_{-\infty}^{\infty} |g(t)|^\beta \right]. \quad (2.83)$$

Dieser Ausdruck (Gleichung (4) in Watson (1979)) ist in der visuellen Psychophysik zum Standardmodell zur Berücksichtigung zeitlicher W-Summation geworden (vergl. auch Watsons Homepage: <http://vision.arc.nasa.gov/mathematica/psychophysica/>).

Ein allgemeiner Ausdruck für Entdecken durch zeitliche W-Summation ist in (2.55), Seite 93, gegeben worden. Die Wahrscheinlichkeit hängt demnach vom Integral der Hazard-Funktion über dem Intervall $[0, T]$ ab. Watson integriert über $(-\infty, \infty)$. Das ist insofern problemlos, als für $t < 0$ $g(t) = 0$ gilt, die untere Grenze also durch 0 ersetzt werden kann. Weiter wird für hinreichend großes t $g(t) = 0$ gelten; Watson nimmt also an, dass der Versuchsdurchgang stets mindestens so lange dauert, bis $g(t) = 0$ gilt. Gleichwohl kann sich hier ein Problem ergeben, denn nach (2.55) hängt die Wahrscheinlichkeit des Entdeckens eben von einer bestimmten Dauer T der Versuchsdurchganges ab, und werden Signale mit verschiedener Dauer betrachtet, werden die Effekte von Darbietungsdauer einerseits und Länge T des Versuchsdurchganges miteinander konfundiert. Weiter ergibt sich die Frage, warum denn überhaupt die Annahme (2.81) gelten soll, d.h. warum von vornherein ein Hochschwellenmodell gelten soll, – der Hinweis auf das Quicksche Modell macht die Annahme verständlich, aber deshalb noch nicht plausibel, zumal der gewissermaßen rasante Übergang zum Integral fragwürdig ist. Zudem impliziert das Hochschwellenmodell, dass die unterliegende zufällige Veränderliche X auf $(-\infty, 0]$ definiert ist, der maximale Wert der repräsentierenden zufälligen Veränderliche also gleich Null ist. Die Frage ist nun, was X bedeutet. Dieser Punkt muß näher erläutert werden.

Extremwertstatistik W-Summation wird für endlich viele stochastisch unabhängige Kanäle bzw. für endlich viele Intervalle mit stochastisch unabhängigen X_j^* zunächst durch (2.62), Seite 96, d.h. durch

$$\psi_n(c) = 1 - \prod_{j=1}^n F_j, \quad 0 < F_j = P(X_j \leq S) < 1$$

charakterisiert; $F_j > 0$ für alle j , weil andernfalls das Produkt stets gleich Null wäre, und $F_j < 1$, weil der Fall $F_j = 1$ das Produkt invariant läßt.

Es sei

$$G_n = \prod_{j=1}^n F_j,$$

so dass $\psi_n(c) = 1 - G_n$, bzw. $G_n = 1 - \psi_n(c)$. Es sei $\psi_n(c) = p$ ein vorgegebener Wert, so dass $G_n = 1 - p$ ebenfalls eine Konstante. Offenbar ist n ein Parameter, von

dessen Wert die Werte der F_j abhängen. Um diese Abhängigkeit zu verdeutlichen, werde zunächst der Einfachheit halber angenommen, dass $F_j = F$ für alle j . Dann gilt

$$\prod_{j=1}^n F_j = F^n = F_n^n \quad 0 < F < 1 \quad (2.84)$$

Die Bezeichnung F_n soll anzeigen, dass F u.a. von n abhängen muß, damit F^n für gegebenen p -Wert den entsprechenden Wert annimmt: denn angenommen, es sei F ein konstanter Wert, $0 < F < 1$; für größere Werte von n würde F^n kleiner werden (für das Produkt $a \cdot b$ zweier Zahlen $a, b \in (0, 1)$ gilt $a \cdot b \leq \min(a, b)$), oder für $n \rightarrow \infty$ würde $F^n \rightarrow 0$ folgen und damit $\psi_n \rightarrow 1$, oder

$$1 - p = F_n^n \Rightarrow \lim_{n \rightarrow \infty} (1 - p)^{1/n} = F_n \rightarrow 1, \quad (2.85)$$

da ja $1/n \rightarrow 0$ für größer werdenden Wert von n . Also müssen die Parameter von $F = F_n$ in bestimmter Weise von n abhängen, damit $F_n = 1 - p$ für verschiedene Werte von n und konstantem Wert für p gilt. Diese Betrachtung überträgt sich auf den Fall unterschiedlicher Verteilungsfunktionen F_j . Man betrachte etwa das geometrische Mittel $\bar{F}$ der F_j :

$$\bar{F} = \left[\prod_{j=1}^n F_j \right]^{1/n}, \quad \text{d.h.} \quad \bar{F}^n = \prod_{j=1}^n F_j.$$

d.h.

$$\psi_n(c) = 1 - \bar{F}^n.$$

Für $\psi_n(c) = p$ und gegebenen Wert von n muß also $\bar{F}$ einen bestimmten Wert haben. Je größer der Wert von n ist, desto größer muß der Wert von $\bar{F}$ sein. So findet man für $p = .75$ und $n = 10$ den Wert $\bar{F} = .871$, für $n = 100$ ergibt sich $\bar{F} = .986$, für $n = 1000$ hat man $\bar{F} = .999$, etc. Je größer der Wert von $\bar{F}$, desto größer müssen auch die Werte der F_j sein. Die für jeden Wert von n richtigen Werte für F_j können demnach mit F_{nj} bezeichnet, um die Abhängigkeit der F_j von n zu betonen. Will man also einen Grenzübergang $n \rightarrow \infty$ durchführen, etwa um zu einer einfachen Integraldarstellung zu gelangen, so muß man dabei berücksichtigen, dass mit wachsendem n die F_{nj} in geeigneter Weise gegen 1 streben müssen, und zwar derart, dass für $n \rightarrow \infty$ das Produkt der F_{nj} immer noch den vorgegebenen Wert $p \neq 0$ hat.

Es sei S ein fester Wert; die Wahrscheinlichkeit, dass die zufällige Veränderliche X_j mit der Verteilungsfunktion F_j einen Wert kleiner (oder größer) als S annimmt, hängt vom Erwartungswert und der Varianz von X_j ab. Man kann deshalb den Ansatz machen, den zu n korrespondierenden Wert von $F_{nj}(S)$ durch eine geeignete Transformation der Skala von X_j zu erreichen: für $X_{nj} = b_n X_j + a_n$ sind der Erwartungswert und die Varianz durch

$$\mathbb{E}(X_{nj}) = b_n \mathbb{E}(X_j) + a_n, \quad \mathbb{V}(X_{nj}) = b_n^2 \mathbb{V}(X_j)$$

gegeben. Die sogenannten *Normierungskonstanten* a_n und b_n müssen so gewählt werden, dass für $n \rightarrow \infty$

$$F_{nj}(S) = P(b_n X_j + a_n \leq S) \rightarrow 1$$

derart, dass

$$\psi_n(c) = \psi(c), \quad 0 < \psi(c) < 1,$$

für alle n . Es läßt sich nun zeigen, dass für eine vorgegebene Verteilung $F_j(S) = P(X_j \leq S)$ die Normierungskonstanten a_n und b_n entweder existieren oder nicht, und dass sie, wenn sie existieren, für die Verteilungsfunktion F_j charakteristisch sind, und dass $P(b_n X_j + a_n)$ für $n \rightarrow \infty$ gegen eine von nur drei möglichen Grenzverteilungen strebt. Eine analoge Aussage gilt für die Verteilung des Minimums von n zufälligen, unabhängig und identisch verteilten Variablen. Die Weibull-Verteilung ist eine dieser drei Grenzverteilungen, für die überdies gilt, dass die Grenzverteilung von Weibull-verteilten Variablen wiederum Weibull-verteilt ist. Eine knappe Einführung in die Theorie der Extremwertstatistik findet man in Mortensen (1998), Kapitel 7; hier soll nur auf einige Ergebnisse eingegangen werden. Führt man nun den Grenzübergang korrekt unter der Berücksichtigung der Re-Skalierung der X_j für Weibull-verteilte Variablen aus, so folgt nicht die Watsonsche Formel (2.83), sondern

$$\psi(c) = 1 - (1 - \gamma) \exp \left[-\frac{1}{T} \int_0^T (g(t, c) - S)^\beta dt \right], \quad g(t) - S \geq 0. \quad (2.86)$$

(Mortensen (2007a)). Hier ist angenommen worden, dass alle X_j^+ – auch im Falle $n \rightarrow 0$, so dass die x_j^+ beliebig nahe benachbart sind, stochastisch unabhängig sind, dass also das Rauschen weiß ist. Für den Fall $c = 0$, der $g(t; 0) = 0$ für alle $t \in [0, T]$ bedeuten soll, folgt dann aber (für $\gamma = 0$)

$$\psi_W(0) = 1 - \exp \left[\frac{1}{T} (-S)^\beta \right] = 1 - \exp \left[-\frac{1}{T} T (-S)^\beta \right] = 1 - \exp \left[-(-S)^\beta \right], \quad (2.87)$$

d.h. die Wahrscheinlichkeit eines echten Falschen Alarms ist unabhängig von der Dauer T des Intervalls, im Gegensatz zu (2.55), wo gezeigt wurde, dass die Wahrscheinlichkeit, dass $\xi(t)$ mindestens einmal während des Intervalls $[0, T]$ den Wert S erreicht, von der Dauer T des Versuchsdurchgangs abhängt. Das gilt natürlich auch für den Spezialfall $S = 0$, d.h. für das Hochschwellenmodell. Man erinnere sich, dass nach (2.55) stets

$$\psi(c) = 1 - \exp \left[-\int_0^T \phi(t, c) dt \right]$$

gilt. $\psi(0) = 0$ impliziert dann

$$\int_0^T \phi(t) dt = 0,$$

und diese Bedingung impliziert wegen $\phi(t) \geq 0$ für alle t den Spezialfall $\phi(t) = 0$ für alle t , d.h. die Dichtefunktion f der Wartezeit $\tau \in [0, T]$, also der Zeit, zu der zum ersten Mal $\xi(t) \geq S$ ist, ist identisch gleich Null.

Ein zu (2.85) analoges Ergebnis findet man, wenn man statt der Weibull-verteilten X_j Gauß-verteilte Variable betrachtet:

$$\psi_G(0) = 1 - \exp(-e^{-S}) \quad (2.88)$$

(Mortensen (2007a), Gl. (14)) Das Ergebnis überträgt sich auf den Ortsbereich, wo dann statt über ein Zeitintervall über ein Ortsintervall integriert wird.

In (2.55), Seite 93, ist die allgemeine Form der psychometrischen Funktion für den Fall der zeitlichen W-Summation angegeben worden; es gilt demnach allgemein

$$\psi(c) = 1 - \exp \left[- \int_0^T \phi(t) dt \right];$$

zeitliche W-Summation kann ja als ein Wartezeitproblem aufgefasst werden. Die Weibull-Verteilung vom Typ (A.182) (Anhang, Seite 194) kann als Wartezeitverteilung interpretiert werden. Die Dichte der Wartezeit ist dann durch (A.184) gegeben, also hat man die Hazard-Funktion

$$\phi(t) = \frac{\beta}{a} \left(\frac{y - \eta_0}{c} \right)^{\beta-1} \exp \left[- \left(\frac{y - \eta_0}{a} \right)^{\beta} \right], \quad y \geq \eta_0 \quad (2.89)$$

Leider ist nicht klar, wie die Systemantwort $g(t)$ in diese Hazard-Funktion eingehen soll, – der Hinweis, dass die Weibull-Verteilung auch eine Verteilung einer Wartezeit ist, ist also so lange nicht hilfreich, wie nicht gesagt wird, wie nun g in ϕ eingeht.

Kritische Anmerkungen Quicks Modell und die diversen Adaptationen dieses Modells – für die zeitliche W-Summation insbesondere Watsons (1979) Modell, für W-Summation zwischen Ortsfrequenzen (Graham (1977), für Kombinationen von Ortsfrequenzen und Positionen im Ortsbereich (Wilson und Bergen (1979), etc. – ist zu einem Standardmodell geworden, dessen Grundannahmen gar nicht mehr hinterfragt werden. Das folgende Zitat ist in vielerlei Hinsicht typisch:

”The probability-summation models suppose that the filtered response to a stimulus is perturbed by additive noise that influences the stimulus visibility (Sachs et al. 1971, Quick 1974, Watson and Nachmias 1977, Watson 1979). Regarding noise fluctuations as independent from point-to-point and moment-to-moment, the probability of seeing the moment is expressed as the joint probability that one sample in space and time of the noisy response has exceeded the criterion level at least once. The probability-summation models for contrast detection are able to explain a lot of data of contrast detection.” Aus Manahilov and Simpson (1999), p. 61.

Diese Bemerkungen sind insofern korrekt, als das Quick-Modell und seine verschiedenen Anwendungen im Zeit- und Ortsbereich in der Tat stets in der genannten Art angewendet wird: es wird nicht hinterfragt, ob die Annahmen (i) vollständige stochastische Unabhängigkeit, und (ii) Gültigkeit der Hochschwellenannahme überhaupt gelten. Plausibel sind diese Annahmen keineswegs. Die generelle Interpretation des Exponenten β als ”Rauschparameter” ist ja nur gerechtfertigt in Bezug auf die Varianz der zufälligen Veränderlichen $\xi(t)$; – hier muß darauf hingewiesen werden, dass

$\xi(t)$ einerseits als Trajektorie eines stochastischen Prozesses betrachtet werden kann, wobei die Schreibweise $\xi(t)$ irreführend ist, denn es ist bei dieser Interpretation ja nicht der Wert der Funktion ξ zum Zeitpunkt t gemeint. Ist dieser Wert gemeint, so wird ξ als zufällige Veränderliche zum Zeitpunkt t aufgefasst, die einen Erwartungswert $\mathbb{E}(\xi(t))$ und eine Varianz $\mathbb{V}(\xi(t))$ hat. β bestimmt die Varianz dieser zufälligen Veränderlichen, – aber nicht die Autokorrelation, von der ja implizit angenommen, dass sie durch die Dirac-Delta-Funktion gegeben ist. Diese Annahme ist aber keineswegs vernünftig und kann zu Fehlinterpretationen der Daten führen.

Farbiges Rauschen In Mortensen (2007a) wird angenommen, dass das Rauschen durch einen Gauß-Prozess mit einer Autokorrelation $R(\tau)$ ist, für die $R(\tau) \neq \sigma^2 \delta(\tau)$ gilt; $R(\tau)$ ist beliebig bis auf die Einschränkung, dass R in eine Taylorreihe entwickelbar sein muß:

$$R(\tau) = R(0) + \tau R'(0) + \frac{\tau^2}{2} R''(0) + 0(\tau), \quad (2.90)$$

wobei $0(\tau)$ eine Funktion von τ ist, die schneller gegen Null geht als τ . In diesem Fall resultiert für die zeitliche W-Summation die psychometrische Funktion

$$\psi(c) = 1 - \exp \left[-\frac{\sqrt{\lambda}}{2\pi} \int_0^T \exp \left(-\frac{(S - g(t; c))^2}{2} \right) \right], \quad 0 < \lambda < \infty \quad (2.91)$$

Hierin ist $\lambda_2 = -R''(0)$ das *zweite Spektralmoment*. λ_2 genügt der Bedingung

$$0 < \lambda_2 < \infty. \quad (2.92)$$

Für kleine Werte von λ_2 fluktuiert das Rauschen langsam, für große schnell: für $\lambda_2 \rightarrow 0$ ist $\xi(t) \approx \xi_0$ eine Konstante, und für $\lambda_2 \rightarrow \infty$ approximieren die Trajektorien $\xi(t)$ das Weiße Rauschen. Der Ausdruck (2.89) wird die Basis für die Diskussion der Frage sein, ob in einem Experiment das Entdecken Peak-Detection oder zeitliche W-Summation ist.

Kapitel 3

Visuelle Psychophysik

3.1 Psychophysik im Zeit- und Ortsbereich

Die Beziehung zwischen Stimulus und Wahrnehmung des Stimulus ist der Gegenstand der Psychophysik, wobei ein Stimulus ein kleiner Lichtfleck oder eine komplette Szenerie, ein einfacher Sinuston, eine Sinfonie oder der Lärm einer Straßenkreuzung sein kann.

Visuelle Szenen können, wenn man von der Farbe absieht, durch eine Funktion der Form $\sigma(x, y, t)$ beschrieben werden, wobei σ ein Luminanzwert ist, (x, y) Ortskoordinaten (retinale Koordinaten) sind und t die Zeit ist. Insbesondere kann ein Stimulismuster in der Form $\sigma(x, y, t) = s(x, y)r(t)$ geschrieben werden. Zur Vereinfachung wird im Folgenden die t -Variable unterdrückt, wenn die zeitliche Modulation nicht explizit diskutiert werden soll, und $s(x, y)$ wird als Luminanzwert am Ort (x, y) genommen.

Das visuelle System kann als ein einziger "Übertragungskanal" oder als "Mehrkanalsystem" konzipiert werden. Ein Kanal kann – für Kontraste im Schwellenbereich – durch eine Impulsantwort $h(x, y, t)$ bzw. durch die Systemantwort $H(u, v, \tau)$ beschrieben werden, wobei H die Fouriertransformierte von h ist. Wird das System im zeitlich eingeschwungenen Zustand betrachtet, kann die Zeitkomponente vernachlässigt werden und man betrachtet das Paar

$$h(x, y) \Leftrightarrow H(u, v), \quad (3.1)$$

wobei

$$H(u, v) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} h(x, y) e^{-i(ux+vy)} dx dy, \quad (3.2)$$

$$h(x, y) = \frac{1}{4\pi^2} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} H(u, v) e^{i(ux+vy)} du dv \quad (3.3)$$

H wird wie folgt interpretiert: nach (A.146) im Anhang ist

$$w = \sqrt{u^2 + v^2} \quad (3.4)$$

die Ortsfrequenzkomponente in der Orientierung (vergl. (A.147))

$$\varphi = \tan^{-1} \frac{v}{u}. \quad (3.5)$$

Die Funktion $h(x, y)$ wird in der deutschen Literatur auch *Punktverwaschungsfunktion* genannt; im Fachjargon hat sich eher der englische Ausdruck *point spread function* (PSF) durchgesetzt. Integriert man h über y , so erhält man die *Linienverwaschungsfunktion*, oder auch *line spread function* (LSF) bezüglich der x -Achse:

$$h_x(x) = \int_{-\infty}^{\infty} h(x, y) dy. \quad (3.6)$$

3.2 Erste empirische Befunde

Im Allgemeinen wird man zunächst einfache Stimulussituationen wählen, um zu grundlegenden Einsichten zu gelangen. Kelly (1972) merkt an, dass allein Experimente mit intermittierendem Licht, also mit periodischen Folgen von Lichtblitzen oder sonstwie zeitlich moduliertem Licht, seit mehr als 200 Jahren durchgeführt werden. Es seien zunächst einige empirische Befunde genannt:

- **Webers Gesetz:** Es sei I die Intensität eines Stimulus, und ΔI der Zuwachs an Intensität, der nötig ist, damit eine Person den Stimulus mit der Intensität I und den mit der Intensität $I + \Delta I$ gerade unterscheiden kann. Nach Weber (1834)¹ soll der Quotient $(I + \Delta I)/I$ gleich einer Konstanten, also unabhängig von I sein. Dieses "Gesetz" gilt nur approximativ in einem bestimmten Bereich von I .
- **Fechners Maßformel:** G. A. Fechner leitete aus dem Weberschen Gesetz seine berühmte Maßformel ab, nach der die Empfindungsstärke proportional zum Logarithmus der Stimulusstärke ist.
- **Bloch's Gesetz (1885):** Es sei T die Dauer, mit der ein Stimulus gezeigt wird, und L_p sei die Luminanz, die notwendig ist, damit der Stimulus in $p\%$ der Darbietungen entdeckt wird. Dann soll $L_p T = c$ eine Konstante sein für $t < T_0$; L_p ist demnach eine Funktion der Dauer T . Die Konstante c hängt von einer Reihe von Faktoren ab: von der Hintergrundhelligkeit, vom Adaptationszustand der Versuchsperson, etc. Dies ist ein *Summationsgesetz*: für einen größeren T -Wert genügt ein kleinerer Luminanzwert, nämlich $L_p = c/T$, damit der Stimulus mit gegebener Wahrscheinlichkeit entdeckt wird. Die durch L_p erzeugte neuronale Aktivität wird gewissermaßen aufsummiert, um die Schwellenaktivität zu erzeugen.
- **Riccis Gesetz (1877):** Dies ist das Pendant zu Blochs Gesetz im Ortsbereich: Demnach gilt $L_p F = c$, c eine Konstante. F ist die Fläche, die der

¹Weber (1834) De pulsu, resorptione, auditu et tactu annotationes anatomicae et physiologicae. Leipzig, 1834

Stimulus überdeckt (gemessen in Sehwinkel). Spätere, genauere Messungen führten auf die als *Piérons Gesetz* (1920), demzufolge $L_p F^{1/3} = c$ gelten soll.

- **Ferry-Porter-Gesetz:** Dieses Gesetz bezieht sich auf die Frequenz f_c , bei der flickerndes Licht zu einem konstanten Licht verschmilzt: nach Ferry (1892) und Porter (1898) soll

$$f_c = a \log L_m + b$$

gelten, wobei L_m die Luminanz ist, bei der das Licht als perfekt konstant wahrgenommen wird. L_m wiederum ist durch das Gesetz von *Talbot und Plateau* (1834) gegeben, demzufolge

$$L_m = \frac{1}{t} \int_0^t L dt,$$

wobei t eine ganze Zahl von vollständigen Zyklen ist.

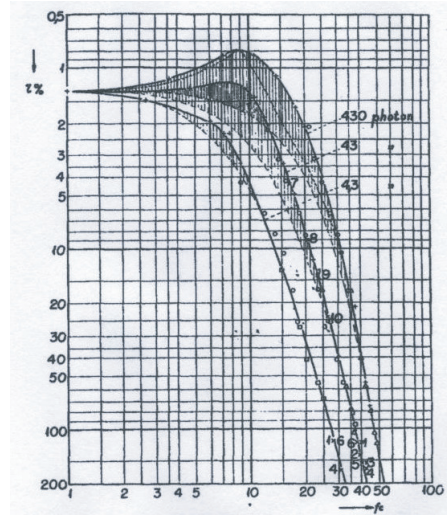
Eine ausführliche Diskussion dieser Befunde findet man in Le Grand (1968). Generell gilt aber, dass es sich bei diesen Gesetzen eigentlich nur um empirische Regelmäßigkeiten handelt, für die der Ausdruck "Gesetz" vielleicht ein wenig zu hoch gegriffen ist, da sie immer nur als Näherungen gelten und die Parameterwerte von den spezifischen experimentellen Bedingungen abhängen. Eine Erklärung dieser Befunde durch Angabe der neuronalen Mechanismen, die die Befunde implizieren, wird zunächst nicht gegeben.

Der systemtheoretische Ansatz in der Psychophysik stellte einen Versuch dar, die empirischen Befunde durch neuronale Modelle zu erklären. Es ist dann zu erwarten, dass durch diesen Ansatz neue Arten von Experimenten entstehen, die sich aus der Konzeption der neuronalen Mechanismen als dynamische Systeme ergeben. Allerdings sagt die Modellierung eines neuronalen Mechanismus noch nichts darüber aus, wie der Entdeckensprozess zu charakterisieren ist, wobei der Begriff des Entdeckens auch das Entdecken von Unterschieden (Diskrimination) umfassen soll. Bevor die Rede auf systemtheoretische Modelle kommt, sollen kurz die üblichen Entdeckensmodelle vorgestellt werden.

3.3 Systemidentifikation im Zeitbereich

Die Frage, bei welcher Frequenz f_c flickerndes Licht zur Wahrnehmung eines homogenen Lichtes in Abhängigkeit von der Hintergrundhelligkeit, dem Adaptationszustand und anderen Variablen führt hat die Forschung spätestens seit der Erfindung der Filmkamera interessiert. f_c ist die *Flimmerverschmelzungsfrequenz* oder *critical flicker frequency* CFF (vergl. das Ferry-Porter Gesetz in Abschnitt 3.2); von dieser Abkürzung wird im Folgenden Gebrauch gemacht. De Lange (1952) scheint der erste gewesen zu sein, der CFFen bestimmt hat, um sie als Frequenzcharakteristik eines linearen, dynamischen Systems zu interpretieren. De Lange konnte nicht mit sinusförmig moduliertem Licht experimentieren, sondern mußte seine periodischen Funktionen durch rotierende Scheiben, auf denen sektorweise weiße und schwarze

Abbildung 3.1: De Langes (1952) CFF-Kurven



Felder angebracht waren, erzeugen. Die Kurven beziehen sich dementsprechend auf die erste Harmonische. Der Abbildung des Logarithmus von c versus den Logarithmus der CFF $f_c(m)$ liefert dann die *Attenuationscharakteristik*² (A.C.), die dem Faktor $A(\omega)$ in Definition 1.2.6, Seite 36, entspricht.

Man kann sagen, dass Flicker dann wahrgenommen wird, wenn die *Maxima* von $g(t)$ einen bestimmten Schwellenwert S erreichen. Sinken sie, bei zu kleinem c , unter diesen Wert, so gilt das Talbot-Plateau-Gesetz und es wird nur eine homogene Helligkeit wahrgenommen. Macht man nun die Peak-Detection-Annahme, so soll also $\max_t g(t) \geq S$ gelten. Die Antwort

$$g(t) = A(\omega) \sin(\omega t + \phi(\omega))$$

hat aber so viele Maxima, wie ωt das Intervall $[0, 2\pi]$ durchläuft.

Abb. 3.1 zeigt zeigt De Langes (1952) Resultate für drei verschiedene Werte der Hintergrundluminanz L_a . Die Abzisse gibt die f_c -, die Ordinate die Attenuation η , d.h. die r -Werte an, wobei r durch

$$r = \frac{\text{Amplitude der Fundamentalfrequenz}}{\text{durchschnittliche retinale Luminanz } L_m} \quad (3.7)$$

definiert ist. De Lange nennt r auch den der *ripple-ratio*. Offenbar haben die A.C.s für größere Hintergrundhelligkeit L_a ein Maximum in der Nachbarschaft von $f_c = 10$ Hz, d.h. die Empfindlichkeit des Systems ist nahe dieser Frequenz am größten. Die Kurven ähneln dann den Resonanzkurven, die man z.B. bei Systemen zweiter Ordnung kennt, und tatsächlich vermutet De Lange (1952), dass es sich hier um Resonanzphänomene handelt. Für niedrige Werte von L_a zeigt das System Tieftaßverhalten. De Lange folgert hieraus, dass das System insgesamt nicht linear ist und

²attenuieren = verdünnen, abschwächen

vermutet, dass der Zusammenhang zwischen der *Form* der A.C.s und der Luminanz L_a sich mathematisch nicht einfach beschreiben läßt.

Die Darstellungen $\log f_c$ versus den Logarithmus der Schwellenmodulation, werden auch kurz *De Lange-Kurven* genannt.

Kellys Daten Für verschiedene Stimuluskonfigurationen ändern sich die De Lange-Kurven nicht qualitativ, aber die genaue Form der Kurven hängt von der Stimuluskonfiguration ab, d.h. von der Größe des Feldes, auf dem der flickernde Stimulus erscheint, und von der Größe des flickernden Stimulus selbst. Bei De Lange hatte das flickernde Feld eine Größe von 2^0 und die stationäre Umgebung eine Größe von 6^0 . Kelly (1959) experimentierte mit einem flickernden 65^0 -Feld, und beobachtete eine leichte Verschiebung der f_c -Werte nach oben und eine *Verringerung* der Empfindlichkeit für *niedrige* Frequenzen. Kellys (1959) Hypothese, dass Augenbewegungen diese Veränderungen der Kurven bewirken, ließ sich nicht bestätigen (Kelly (1972)). Ausgedehnte Untersuchungen mit verschiedenen Reizmustern führten Kelly (1969) zu der Hypothese, dass eine Vergrößerung des Reizfeldes die Empfindlichkeit für *alle* Frequenzen erhöht, dass aber für Frequenzen größer als 10 Hz die *laterale* Hemmung im visuellen System unwirksam wird.

Kellys Untersuchungen zeigten, dass die Kurven $\log c$ versus $\log f[\text{Hz}]$, die für verschiedene L_a -Werte erhoben werden, für hohe Frequenzen eine gemeinsame Asymptote haben, wenn man Luminanzwerte in Trolands ausdrückt, d.h. in Einheiten absoluter retinaler Helligkeit (1 Troland [td] = Luminanz von 1 candela/ m^2 , gesehen durch eine 1 mm^2 -Pupille (zit. nach LeGrand (1968), p. 107)). Für einen großen Bereich von Hintergrundhelligkeiten ($.06 \text{ td} \leq L_a[\text{td}] \leq 9300 \text{ td}$) zeigt sich nun, dass die CFF f_c eine Modulation $m = m/L_a = 1$ gegeben ist; daraus folgt ,

- (i) dass die Antwort auf den flickernden Teil des Stimulus (d.h. auf $\sin(\omega t)$) unabhängig von L_a ist und damit die Annahme gerechtfertigt erscheint, dass das System für die betrachteten Amplituden tatsächlich im linearisierten Bereich arbeitet (Kelly (1961)), und
- (ii) das Ferry-Porter Gesetz modifiziert werden muß. Denn $c = c/L_a = 1$ für alle L_a bedeutet ja, dass die *Modulation* unabhängig von L_a ist. Dagegen hängt c_a , die absolute Amplitude, von L_a ab, und Kelly schlägt statt (??) die Modifikation

$$m(f_c) = a \log L_a + b, \quad (3.8)$$

m für Modulation, vor. Für die Interpretation der Kurven hinsichtlich der Signalübertragungsmechanismen wird (3.8) eine Rolle spielen.

Die Kellyschen Untersuchungen sind hier keineswegs vollständig referriert worden, sondern nur so weit, wie sie für die im folgenden betrachteten Modellbildungen von Bedeutung sind.

Roufs Daten In einer Reihe von Untersuchungen hat Roufs (1972a), (1972b), (1973), (1974) nicht nur CFFen in Abhängigkeit vom jeweiligen Adaptationszustand

untersucht, sondern ebenfalls die Beziehung der entsprechenden De Lange-Kurven zu den Schwellenwerten für Rechteckpulse (vergl. (??)) und für Folgen solcher Pulse. Dabei hat der Reiz geometrisch die Form einer Kreisscheibe mit 1^0 Durchmesser; die Versuchsperson sah den Stimulus durch eine künstliche Pupille von 2 mm Durchmesser.

Die Kurven für die CFFen entsprechen qualitativ den von De Lange und Kelly gewonnenen Kurven, und die Schwellenwerte für Rechtecke zeigen, dass Blochs Gesetz (??) erfüllt ist (Roufs, 1972a). Aus den Schwellenwerten für Rechteckpulse lässt sich der Wert der kritischen Zeit t_b berechnen: es sei ϵ_a der Schwellenwert für Rechteckstimuli mit Dauern $t_s \gg t_b$, für die also keinerlei Summation mehr zu beobachten ist; ϵ_a hängt von L_a ab, und natürlich $\epsilon_a = \textit{konstant}$ für $t_s \gg t_b$. Dementsprechend entspricht den Werten von $t_s \gg t_b$ in Fig. 3.2. x eine Gerade parallel zur x -Achse. Den Zeiten $t_s < t_b$ entspricht die Gerade mit der Steigung -1 . Der Schnittpunkt dieser beiden Geraden kann als Schätzung für t_b gewählt werden. Definiert man $F_a = 1/\epsilon_a$ als "Empfindlichkeit" des Systems (Roufs (1972a)), so zeigen die Roufschen Messungen, dass F_a mit steigendem Wert von L_a kleiner werden, d.h. ϵ_a steigt mit L_a . Dieses Resultat würde man natürlich schon aufgrund der Weberschen Beziehung erwarten. Gleichzeitig fällt aber auch die kritische Zeit t_b mit L_a von 110 ms ($L_a \approx .8 \text{ tr}$) auf 20 ms (10^3 tr).

Während aber die Daten für Darbietungszeiten $t_s < t_b$ stets auf die Gültigkeit des Blochschen Gesetzes weisen, sind die Daten für Zeiten $t_s > t_b$ keineswegs eindeutig. Insbesondere präsentierte Roufs (1974) Daten zur zeitlichen Summation, die weder auf eine Unabhängigkeit der Entdeckungsschwelle für Zeiten größer als t_b noch auf die Beziehung (??) (Piéron's Gesetz) verweisen. Vielmehr findet sich, wenn man nur den Bereich der t_s -Werte in der Nachbarschaft von t_b hinreichend genau untersucht, eine "Delle", oder, wie Roufs sagt, ein "Dip": die Schwellen sinken bis zu einem bestimmten Wert von t_s , um dann wieder anzusteigen, bis ein bestimmtes Niveau erreicht wird, auf dem der Schwellenwert dann bleibt. Für die Modellierung der Dynamik im Zeitbereich hat dieser *Dip* einige Implikationen, auf die in den folgenden Kapiteln noch weiter eingegangen wird.

Roufs (1974) liefert weitere Daten über die Entdeckung von Folgen zweier Rechteckreize ("Doublets"), die u. a. relevant sind für die Frage, ob die Schwellen für Inkremente andere Werte haben als die für Dekremente. Diese Frage ist relevant für die Diskussion der Eigenschaften der Teilsysteme, die von den ON- bzw. den OFF-Zellen gebildet werden (dies ist das B- und D-System Jungs (197?)), bezüglich ihrer Rolle im Mustererkennungsprozess. Boynton, Ikeda und Stiles (1954) fanden für Dekremente einen geringeren Schwellenwert als für Inkremente (fovealer, roter Stimulus mit $10'$ Durchmesser. Patel und Jones (1968) bestätigten das Resultat für einen peripher dargebotenen (7^0) Stimulus, allerdings in Abhängigkeit von der Hintergrundhelligkeit L_a , der Darbietungsdauer sowie vom Durchmesser. Short (1966) führte Schwellenbestimmungen für Stimuli durch, die ebenfalls im Bereich peripheren Sehens (15^0) dargeboten wurden. Er fand, dass die Dekrementschwellen für einen Hintergrund im skotopischen Bereich $.4 \text{ log-Einheiten}$ niedriger waren als die Inkrementschwellen, dass die Unterschiede aber im photopischen Bereich fast verschwanden.

Andere Autoren (Blackwell (1946), Vos, Lazet und Bouman (1956), Herrick (1956) und Rashbass (1970) dagegen fanden gelegentlich kleine, aber im allgemeinen nicht signifikante Differenzen zwischen den beiden Schwellenarten. Roufs (1974) berichtet ebenfalls gelegentliche, aber nicht systematisch auftretende Differenzen.

Es wird zunächst der grundsätzliche Ansatz skizziert. Der Stimulus $s(t)$ erscheine auf einem (Bild-)Schirm mit durchschnittlicher Luminanz L_0 . Durch Präsentation von s wird diese Luminanz moduliert. s spreche einen linearisierbaren Kanal an, der für hinreichend schwache Intensitäten durch eine zeitliche Impulsantwort $h(t)$ charakterisiert ist. Der retinale Bereich, der durch s überdeckt wird, wird dann in der Form

$$L(t) = L_0 + cs(t) \quad (3.9)$$

moduliert, wobei c die Modulation ist: s hat die Modulation von 1, und

$$c = \frac{L_{\max} - L_{\min}}{L_{\max} + L_{\min}}. \quad (3.10)$$

$s(t)$ habe die Fouriertransformierte $S(\omega)$ – er Einfachheit halber wird im Folgenden einfach halber $j = \sqrt{-1}$ fortgelassen –

$$S(\omega) = \int_{-\infty}^{\infty} s(t)e^{-j\omega t} dt, \quad \omega 2\pi f, \quad (3.11)$$

f die Frequenz in Hertz [Hz]. Die Antwort des Kanals ist dann durch die Faltung von s und h gegeben:

$$g(t) = c \int_0^t s(\tau)h(t - \tau)d\tau = \frac{c}{2\pi} \int_{-\infty}^{\infty} H(\omega)S(\omega)e^{j\omega t} d\omega. \quad (3.12)$$

Diese Darstellung der Antwort g gilt nur für kleine Werte von c , und 'klein' wird nun so interpretiert, dass die Antwort im Schwellenbereich ist. Dies soll insbesondere bedeuten, dass die Antwort g nur mit einer bestimmten Wahrscheinlichkeit zu einer Wahrnehmung führt. Es wird nun postuliert, dass die Versuchsperson nach dem Peak-Detection-Kriterium entscheidet: sie signalisiert "ja" (für: es wurde ein Stimulus präsentiert), wenn $\max_t g(t) = g_0(p)$, wobei $g_0(p)$ derjenige Wert von $\max_t g(t)$ ist, für den die Wahrscheinlichkeit des Entdeckens gerade gleich p ist, – etwa $p = .5$ oder $p = .75$.

3.3.1 Schätzung der Amplitudencharakteristik

Das Ziel ist, die Impulsantwort h zu bestimmen. Da die Systemfunktion H die Fouriertransformierte von h ist, kann ebensogut H bestimmt werden. Durch inverse Transformation erhält man dann daraus eine Schätzung für h . Dazu werde

$$s(t) = \sin(\omega_0 t), \quad \omega_0 = 2\pi f_0 \quad (3.13)$$

gewählt; f_0 ist die Frequenz, gemessen in Hertz ([Hz]). Die Fouriertransformierte von s ist dann

$$S(\omega_0) = 2\pi j[\delta(\omega + \omega_0) - \delta(\omega - \omega_0)], \quad \omega = 2\pi f. \quad (3.14)$$

Nach Satz 1.2.10, Seite 37, gilt dann

$$g(t) = c|H(j\omega)| \sin(\omega t + \varphi(\omega)), \quad (3.15)$$

(vergl. (1.82), wobei L in Satz 1.2.10 die Bedeutung eines linearen Operators und nicht, wie hier, der Luminanz hat). Es ist aber

$$\max_t g(t) = g_0(c) = c|H(\omega)| \max_t \sin(\omega t + \varphi(\omega)) = c|H(j\omega)|, \quad (3.16)$$

denn $\max_t \sin(\omega t + \varphi(\omega)) = 1$. Also ist

$$\frac{1}{c} = |H(2\pi f_0)|, \quad f_0 = 2\pi/\omega_0, \quad c = c(f_0), \quad (3.17)$$

wobei c der Kontrast ist, bei dem das Flickern gerade verschwindet (Flickerkontrast). Die Schreibweise $c(f_0)$ soll andeuten, dass c von der Frequenz f_0 der Sinusfunktion abhängt.

Man kann nun $c(f_0)$ für eine Reihe von f_0 -Werten bestimmen und die $1/c(f_0)$ gegen die f_0 bzw. die $\log f_0$ auftragen und erhält so eine Schätzung der Transferfunktion bzw. der Amplitudencharakteristik der Systemfunktion. Macht man zusätzlich die Annahme, dass der getestete Kanal als Minimumphasensystem beschreibbar ist, läßt sich daraus die Systemfunktion selbst abschätzen. Abb. 3.2 zeigt die entsprechend dem Ansatz (3.17) geschätzten Amplitudenspektren für verschiedene Hintergrundhelligkeiten. Die Daten legen ein Bandpassverhalten des entdeckenden Filters nahe, das mit niedriger werdender Hintergrundhelligkeit in ein Tiefpassverhalten übergeht. Die Parameter der Linearisierung hängen also vom Adaptationszustand relativ zur Hintergrundhelligkeit ab.

Rechteckstimuli

Hat man eine Schätzung von $H(\omega)$ gefunden, so läßt sich daraus die Impulsantwort $h(t)$ durch inverse Transformation bestimmen. Man kann dann versuchen, die Schwellendaten für Versuche mit Rechteckpulsen vorherzusagen, – man erhält damit eine Kreuzvalidierung der Ergebnisse. Rechteckpulse sind gemäß

$$s(t) = \begin{cases} 1, & 0 \leq t \leq T_s \\ 0, & t < 0 \text{ oder } t > T_s \end{cases} \quad (3.18)$$

definiert. Die Antwort den entdeckenden Kanals mit der Impulsantwort h ist

$$g(t) = \int_0^{t_0} h(t_0 - \tau) s(\tau) d\tau, \quad t_0 = \begin{cases} t, & t_0 \leq T_s \\ T_s, & t > T_s \end{cases} \quad (3.19)$$

Eine generelle Frage ist, ob Befunde wie Blochs Gesetz, also $c(T_s)T_s = k$ eine Konstante, erklärt werden können, wobei $c(T_s)$ die Schwellenintensität eines Rechteckpulses der Dauer T_s ist, und $T_s < T_0$, wobei T_0 eine kritische Dauer des Pulses ist: bis zu dieser Dauer soll das Blochsche Gesetz gelten, darüber hinaus nicht mehr.

Abbildung 3.2: (a) Sinusoidaler Flickerstimulus und vermutete Form des Amplitudenspektrums. (b) Schätzungen des Amplitudenspektrums für verschiedene Werte der Hintergrundshelligkeit (gemessen in Troland). Die Kurven legen für niedrige Hintergrundshelligkeiten ein Tiefpass-, für höhere Hintergrundshelligkeiten ein Bandpassverhalten des entdeckenden Filters nahe. Aus Roufs (1972a).

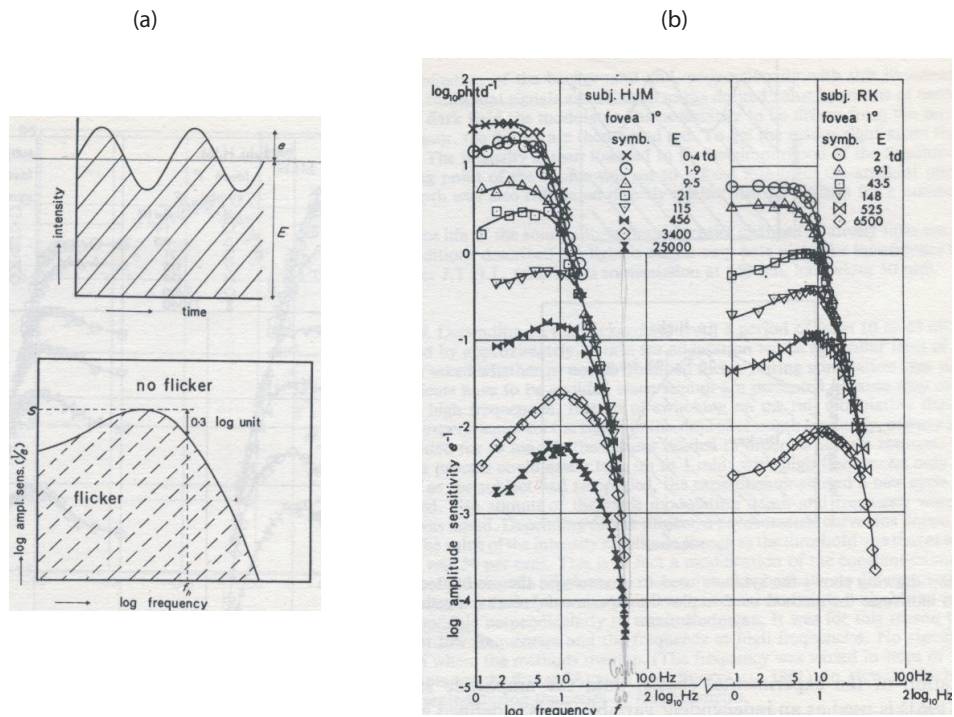
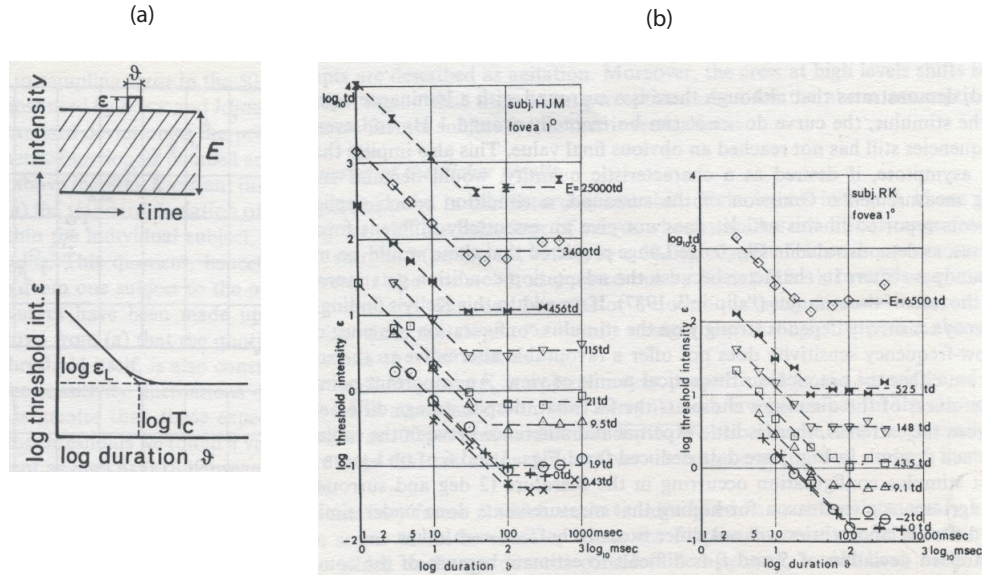


Abbildung 3.3: (a) Schema des Rechteckpulses und Vorhersage des Blochschen Gesetzes. ε ist die Intensität des Pulses, $\vartheta (= T_s)$ seine Dauer. (b) Daten für verschiedene Hintergrundshelligkeiten. Aus Roufs (1972a).



In Abb. 3.3 wird in (a) das Schema der Vorhersage für das Blochsche Gesetz gezeigt (nach Roufs (1972a)): logarithmiert man das Produkt $c(T_s)T_s = k$, so erhält man

$$\log c(T_s) + \log T_s = \log k. \quad (3.20)$$

$\log c(T_s)$ ist demnach eine lineare Funktion von $\log T_s$, und jenseits einer kritischen Dauer bleibt $\log c(T_s)$ konstant. Die in (b) gezeigten Daten (für verschiedene Hintergrundshelligkeiten) sind mit dieser Vorhersage kompatibel. Wie in Abb. 3.2, so sind die Daten abhängig vom Hintergrund, – was natürlich zu erwarten ist.

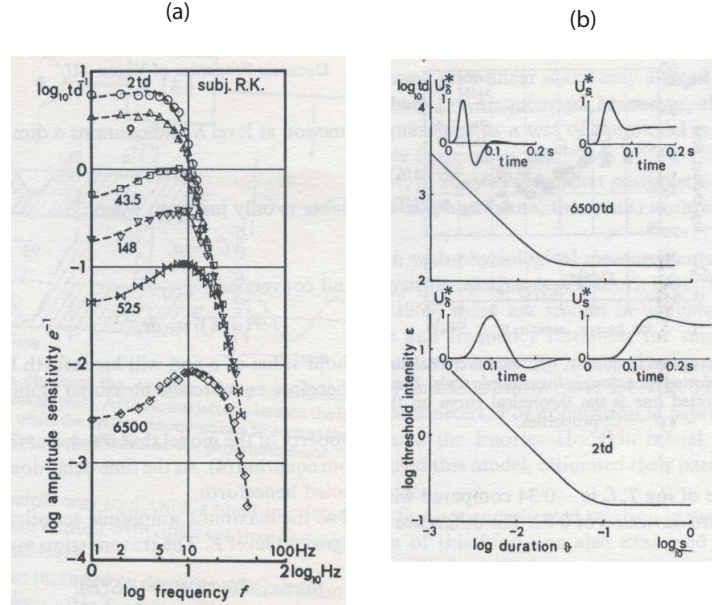
Man wird die Frage nur im Zusammenhang mit einer Annahme über den Entdeckensmechanismus diskutieren können. Wie bei der Interpretation der sinusoidalen Flickerstimuli liegt es nahe, zunächst den Peak-Detection-Mechanismus anzunehmen. Dann wird man auf die Vermutung geführt, dass das Blochsche Gesetz vermutlich eine Implikation der Form von g in Abhängigkeit von der Dauer T_s des Stimuli ist. Dazu werde zunächst angenommen, dass $T_s = \infty$ ist. Praktisch bedeutet dies, dass T_s gleich der Dauer des Versuchsdurchgangs ist, dass also der Stimulus etwa eine Sekunde gezeigt wird, oder vielleicht zwei. Der Stimulus wird dann durch eine Schrittfunktion definiert:

$$s(t) = \begin{cases} 1, & t \geq 0 \\ 0, & t < 0 \end{cases}, \quad (3.21)$$

und g ist dann durch

$$g(t) = c \int_0^t h(t - \tau) s(\tau) d\tau \quad (3.22)$$

Abbildung 3.4: (a) Amplitudenspektren mit angepasster Vorhersage durch ein Filtermodell. (b) Aus dem Modell für die Amplitudenspektren vorhergesagte Impuls- und Schrittantworten. Ein expliziter Ausdruck für die Impulsantworten wird nicht angegeben, die Funktionsverläufe der Impulsantworten ergeben sich durch numerische Inversion der Transferfunktion (3.23). Aus Roufs (1972b).



erklärt. Es werde nun angenommen, dass das Integral mit t monoton wächst. Dann ist

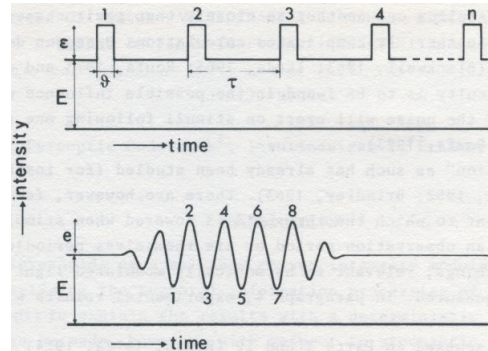
$$\max_t g(t) = g(t)$$

und damit $c(T_s)g_0(T_s) = k$ eine Konstante gilt, muß $c(T_s)$ monoton mit T_s fallen, im Widerspruch zum Blochschen Gesetz, nach dem der Trade-off zwischen T_s und $c(T_s)$ für Werte $T_s > T_0$ nicht mehr gilt. Dies bedeutet, dass h nicht von der Art ist, dass das Integral monoton mit T_s wächst, sondern für ein T_0 ein Maximum annimmt und danach kleiner oder höchstens gleich diesem Maximum ist. Solche Impulsantworten wird man für Kanäle vermuten, in denen Rückkopplungsmechanismen wirken, die einen inhibitorischen Effekt haben. In Abb. 3.4 werden noch einmal die Amplitudenspektren gezeigt, zusammen mit den Vorhersagen eines Filtermodells. Es wurde angenommen (Roufs (1972b)), dass der entdeckende Filter ein Minimum-Phasen-System ist, das durch die Transferfunktion

$$H(\omega) = C_f \frac{G_t^{10}}{(1 + j\omega T_i)^{10}} \frac{(G_d + j\omega T_d)^2}{(1 + j\omega T_d)^2} \quad (3.23)$$

charakterisiert ist, wobei G_i und G_d bestimmte Konstante (*Gain factors*) sind, und T_i und T_d Zeitkonstanten. C_f ist eine Photometerkonstante. Die inverse Transformation von H liefert die Impulsantworten in Abb. 3.4. Offenbar besteht der Filter

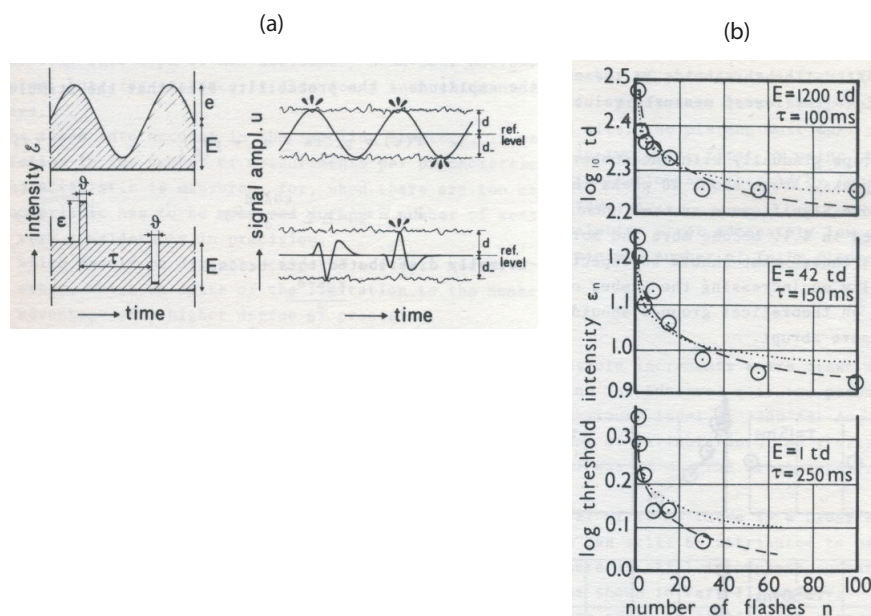
Abbildung 3.5: Flash-trains und "gated sinuoids"



insgesamt aus 12 hintereinandergeschalteten Teilfiltern, von denen jeweils 10 bzw. 2 von gleicher Bauart sind. Die korrespondierenden Impulsantworten (durch numerische Inversion von H gewonnen) zeigen deutlich die inhibitorische Reaktion, die nach einer von der Hintergrundhelligkeit abhängenden Zeit eintritt und die Eindeutigkeit von $\max_t g(t)$ bedingen. Roufs (1972b) merkt an, dass andere, ähnlich definierte Transferfunktionen die Daten ebenfalls gut beschreiben. Hier wird eine Problematik der Charakterisierung neuronaler Systeme durch psychophysische Daten der betrachteten Art deutlich: die mangelnde Eindeutigkeit. Natürlich ist diese Problematik typisch für jeden Versuch, empirische Daten durch den Fit von Modellen zu erklären. Einzelne Datensätze werden kaum jemals nur durch ein Modell zu erklären sein. Es soll deshalb auch nicht weiter versucht werden, die neuronalen Mechanismen zu diskutieren, die der Transferfunktion (3.23) entsprechen. Der Hauptzweck der Präsentation der Roufschen Daten ist denn auch, das Prinzip zu illustrieren, nach dem vorgegangen werden kann, empirische Befunde systemtheoretisch zu deuten.

Die Beziehung zwischen Flicker- und Pulsdaten wurde hergestellt durch die Annahmen (i) linearisierbarer Filter, und (ii) Entdecken durch Peak-Detection. Die Vorhersage des Blochschen Gesetzes scheint diese beiden Annahmen zu rechtfertigen, allerdings sind nicht alle Daten, die in der Literatur zu finden sind, so eindeutig. Die Amplitudenspektren variieren nicht nur mit der Hintergrundhelligkeit, sondern auch mit der spatialen Struktur des Stimulus, und die Möglichkeit der Wahrscheinlichkeitssummation kann insofern nicht ausgeschlossen werden, als z.B. die Darbietungsdauer des flickernden Stimulus selten explizit kontrolliert wurde. So ist es denkbar, dass bei der Darbietung kurzer zeitlicher Rechteckstimuli die Peak-Detection-Annahme gilt, bei der Darbietung von Sinusflicker aber dennoch Wahrscheinlichkeitssummationseffekte auftreten. Denn jedesmal, wenn $c \sin(2\pi ft)$ maximal, also gleich 1, wird, wird ja (nach von der Frequenz f abhängender Phasenverschiebung ein "Peak" der internen Antwort erzeugt. Nimmt man etwa an, dass es eine langsam variierende stochastische Komponente in der Aktivierung gibt, so ist denkbar, dass einige dieser Maxima den Wert der internen Schwelle erreichen, andere aber

Abbildung 3.6: Flash-trains und "gated sinusoids" – (a) Schema, (b) Daten



nicht. Je länger der flickernde Stimulus präsentiert wird, desto größer wird dann die Wahrscheinlichkeit, der einer der Maxima den Schwellenwert erreicht und die V_p eine "Ja"-Antwort gibt. Dieser Wahrscheinlichkeitssummationseffekt würde niedrigere Amplitudenschätzungen erzeugen als ein "reiner" Peak-Detection-Mechanismus.

Stochastische Effekte und Peak-Detection Die vorangegangenen Überlegungen legen nahe, dass man die Peak-Detection-Annahme nicht unhinterfragt übernehmen kann. In der Tat hat Roufs (1974) einen ersten Test unternommen. Dazu zeigte er "flash trains", d.h. Folgen von Rechteckpulsen, und zeitlich sinusförmig modulierte Stimuli, vergl. Abb. 3.5.

In Abb. 3.6 (a) wird der Entdeckensprozess angedeutet: der Grundmechanismus ist Peak-Detection, da es aber Folgen von Peaks gibt, kann es zwischen den Peaks zu Wahrscheinlichkeitssummationseffekten kommen. Die Anzahl n der Rechteckpulse kann systematisch variiert und für jeden Wert von n kann die Schwellenintensität bestimmt werden, für die der "train" gerade wahrgenommen wird. Das gleiche kann für "gated sinusoids" durchgeführt werden. Hier wird der Sinusflicker nicht mit voller Amplitude eingeschaltet, sondern die Amplitude wird, wie in der Abb. 3.5 angedeutet, langsam erhöht. Dadurch werden transiente Effekte in der Antwort g vermieden. Variiert wird die Anzahl der Zyklen für feste Frequenz f . Abbildung 3.6 (b) zeigt die resultierenden Daten. Mit größer werdendem Wert von n fällt die Schwellenintensität bzw. – kontrast ab. Der Abfall entspricht der Annahme, dass die jeweiligen Maxima unabhängig voneinander zur Entdeckung des Gesamtstimulus beitragen. Die Ergebnisse zeigen, dass Peak-Detection entweder nicht streng gilt, oder abhängig von der

Instruktion ist. In jedem Fall legen die Befunde von Roufs (1974) nahe, dass man nicht unbenommen Flickerdaten, aus denen unter Annahme des Peak-Detection-Mechanismus Schätzungen der System- bzw. Transferfunktionen gewonnen wurden, nicht notwendig valide Schätzungen der Impulsantworten implizieren. Es folgt, dass direkte Schätzungen der Impulsantworten notwendig sind. Darauf wird im folgenden Abschnitt eingegangen.

3.3.2 Die Perturbationsmethode

Diese Methode wurde von Roufs und Blommaert (1981) vorgeschlagen. Das Ziel der Experimente, bei denen diese Methode angewendet wird, ist die direkte Bestimmung des zeitlichen Verlauf der Antwort auf einen Stimulus, dessen Verlauf im Prinzip beliebig sein darf. Das Wort "direkt" darf nicht allzu wörtlich genommen werden, wie die folgende Beschreibung der Methode deutlich machen wird; gemeint ist die "indirekte", über die Schätzung der Transferfunktion gewonnene Abschätzung der Impulsantwort.

Die Methode basiert einerseits auf dem Superpositionsprinzip, das für beliebige lineare Systeme gilt, andererseits auf der Annahme des Peak-Detection-Prinzips. Gelten beide Annahmen, so sollten aus den Daten, die anhand sehr kurzer Rechteckpulse erhoben wurden, die Daten für einen als Schrittfunktion präsentierten Reiz vorhergesagt werden können und umgekehrt, aus den Daten für die Schrittfunktion sollten die Impulsfunktion vorhergesagt werden können. Diese Kreuzvorhersagen liefern einen Test für beide Annahmen: Linearität und Peak-Detection.

Es sei $s(t)$ die Funktion der Zeit, die den zeitlichen Verlauf des Stimulus beschreibt, der bestimmt werden soll; in Abbildung 3.7 ist dies ein Rechteckpuls. Er wird auf einer Hintergrundluminanz E präsentiert. In Roufs Blommaerts Notation ist die Intensität dieses Stimulus ε . Der perturbierende Stimulus s_p ist stets ein sehr kurzer Puls der Dauer θ :

$$s_p(t) = \begin{cases} 1, & 0 \leq t \leq \theta \\ 0, & t < 0, \quad t > \theta \end{cases} \quad (3.24)$$

Für hinreichend kleinen Wert von θ kann man dann $s_p(t) = \theta \delta(t)$ schreiben, so dass die Antwort auf s_p durch

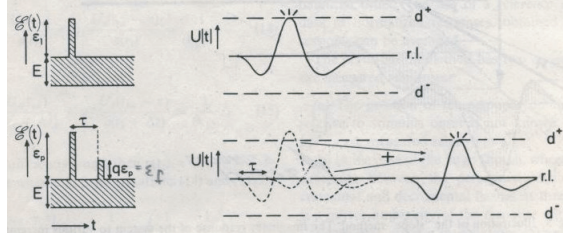
$$g_p(t) = c_p \theta \int_0^t \delta(t') h(t - t') dt' = c_p \theta h(t)$$

gegeben ist³. Für hinreichend kleinen Wert von θ ist also

$$u_p(t) = c_p \theta h(t). \quad (3.25)$$

³Das kann auch explizit gezeigt werden: es sei $H(t) = \int h(\tau) d\tau + C$, C eine Integrationskonstante. Dann gilt insbesondere $\int_0^\theta h(t - \tau) d\tau = \int_{t-\theta}^t h(t') dt' = H(t) - H(t - \theta)$, und wegen $\lim_{\theta \rightarrow 0} \frac{H(t) - H(t - \theta)}{\theta} = h(t)$ folgt für hinreichend kleinen Wert von θ dann $H(t) - H(t - \theta) \approx \theta h(t)$, so dass (3.25) folgt.

Abbildung 3.7: Das Prinzip der Perturbationsmethode (Roufs und Blommart (1981))



Es sei weiter s die Funktion der Zeit, die den zeitlichen Verlauf eines Stimulus definiert; es soll der zeitliche Verlauf der entsprechenden Antwort des entdeckenden neuronalen Kanals bestimmt werden. s wird zeitlich versetzt zum Puls s_p präsentiert, d.h. es wird die Superposition

$$\sigma(t) = c[s_p(t) + qs(t - \tau)]$$

gezeigt; der Puls hat also die Intensität c , und s hat die Intensität qc . Die Antwort des Kanals auf $s(t - \tau)$ sei $u(t - \tau)$. Gilt die Peak-Detection-Annahme, so wird der Stimulus entdeckt, wenn

$$\max_t [c(\theta h(t) + qu(t - \tau))] = k, \quad (3.26)$$

k eine bestimmte Konstante. Es seien

$$h'(t) = \frac{dh(t)}{dt}, \quad u'(t - \tau) = \frac{du(t - \tau)}{dt}.$$

Mit $\sigma(t) = \theta h(t) + qu(t - \tau)$ gelte $\sigma(t_{\max}) = \max_t \sigma(t)$, d.h. es gilt dann

$$\left. \frac{d\sigma(t)}{dt} \right|_{t=t_{\max}} = c[\theta h'(t_{\max}) + qu'(t_{\max} - \tau)] = 0,$$

so dass

$$h'(t_{\max}) = -qu'(t_{\max} - \tau). \quad (3.27)$$

Das Maximum der Impulsantwort h werde zum Zeitpunkt t_{ex} erreicht, so dass $h'(t_{ex}) = 0$. q werde nun hinreichend klein gewählt, so dass $qu'(t_{\max} - \tau) \approx 0$, d.h. $h'(t_{\max}) \approx 0$. Für hinreichend kleinen Wert von q kann dann

$$t_{\max} \approx t_{ex} \quad (3.28)$$

vermutet werden. Aus der Peak-detection-Annahme und (3.26) folgt einerseits

$$\theta h(t_{\max}) + qu(t_{\max} - \tau) = \frac{k}{c(\tau)}, \quad (3.29)$$

sowie

$$\theta h(t_{ex}) = \frac{k}{c_p}, \quad (3.30)$$

Abbildung 3.8: Impuls- (a) und Schrittfunktionsantwort (b) für einen phasischen Kanal (Roufs und Blommart (1981))

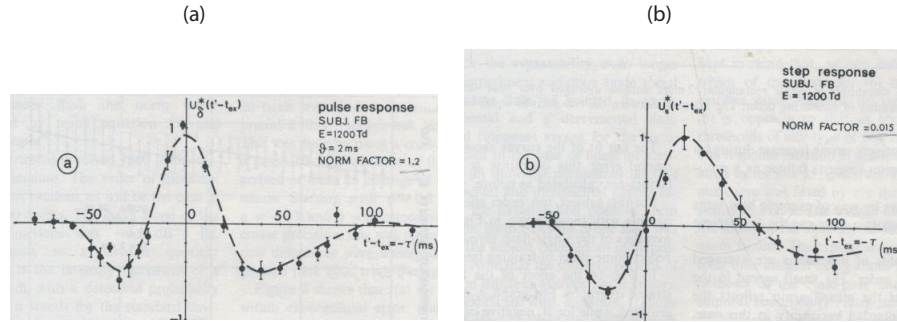
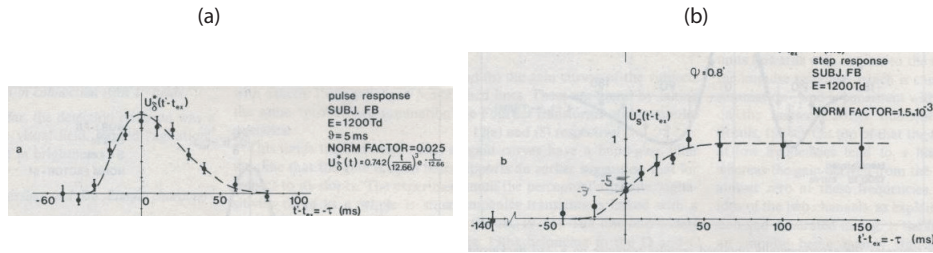


Abbildung 3.9: Impuls- (a) und Schrittfunktionsantwort (b) für einen tonischen Kanal (Roufs und Blommart (1981))



wobei $c(\tau)$ die Schwellenintensität der Superposition und c_p die Schwellenintensität des Pulses ist, wenn er allein, also ohne s gezeigt wird. Wegen (3.28) wird nun $\theta h(t_{\max})$ in (3.29) durch k/c_p ersetzt und man erhält

$$\frac{k}{c_p} + q u(t_{\max} - \tau) = \frac{k}{c(\tau)},$$

woraus man die Schätzung

$$\frac{1}{k} u(t_{ex} - \tau) \approx \frac{1}{q} \left(\frac{1}{c(\tau)} - \frac{1}{c_p} \right), \quad -t_{ex} \leq \tau \quad (3.31)$$

erhält. Trägt man nun die rechte Seite gegen τ auf, so erhält man eine – approximative – Darstellung von $u(t_{ex} - \tau)/k$. Die Rechtfertigung, $t_{\max}$ durch t_{ex} zu ersetzen liegt in (3.28) sowie in der Notwendigkeit, eine definitive Zeitskala zu haben, in Bezug auf die u dargestellt werden kann, denn der exakte Wert von $t_{\max}$ hängt von τ ab, wobei in (3.31) eben angenommen wird, dass diese Abhängigkeit vernachlässigbar ist. In Abb. 3.8 werden die Verläufe der Impulsantwort und der zugehörigen Schrittantwort für einen phasischen Kanal gezeigt. Die Schrittantwort kann durch

(numerische) Integration der Impulsantwort nachgerade perfekt vorausgesagt werden; dies stützt die Annahmen der Linearität und des Peak-Detection, und natürlich der Annahmen, die der Perturbationsmethode unterliegen, insbesondere die Annahme (3.28), – oder der Annahme, dass $h(t_{ex}) \approx h(t_{\max})$, denn dies ist die wesentliche Annahme, da $1/c_p$ für $h(t_{ex})/k$ substituiert wird. Der Aspekt der Impulsantwort, der Widerspruch evoziert hat, ist die inhibitorische "Delle" (von - 50 ms bis - 20 ms), die bei den inversen Transformationen von Amplitudenspektren auf der Basis der Annahme der Minimalphaseneigenschaft nicht gefunden wurde. Watson (1983) hat argumentiert, dass diese Delle ein Artefakt ist, das durch Nichtbeachtung der zeitlichen Wahrscheinlichkeitssummationseffekte zustande käme. Auf seine Alternative, die anhand einer speziellen

3.3.3 Peak-Detection versus Wahrscheinlichkeitssummation

Der Perturbationsmethode nimmt außer der Linearität des entdeckenden Kanals noch an, dass (i) das Entdecken als Peak-detection charakterisiert werden kann, und (ii) dass $t_{ex} \approx t_{\max}$ gilt (vergl. (3.28). Roufs und Blommaert (1980) hatten versäumt, explizit darauf hinzuweisen, dass $t_{\max}$ im Allgemeinen nur in der Nachbarschaft von t_{ex} liegen wird und haben statt dessen einfach die Gleichheit dieser beiden Zeitpunkte postuliert⁴. Dies brachte ihnen die herbe Kritik Watsons (1981) ein, der den inhibitorischen ersten Teil der Impulsantwort für phasische Kanäle für ein Artefakt der Perturbationsmethode hielt, das im allerdings im Wesentlichen auf die seiner Ansicht nach nicht zu rechtfertigende Annahme der Peak-Detection zurückzuführen sei. Denn die plausible Annahme sei Entdecken durch zeitliche Wahrscheinlichkeitssummation, das er durch seinen Ansatz (??), Seite ?? erklären zu können meinte. In einem ersten Schritt der Kritik schlug Watson (1981) eine Impulsantwort, deren Fourier- bzw. Laplacetransformierte ein Amplitudenspektrum lieferte, das exakt einem von Roufs und Blommaert (1980) erhobenen Amplitudenspektrum für den gleichen Stimulus entsprach. Die allgemeine Form dieser Impulsantwort ist

$$h_w(t) = (at)^{p_1} e^{-at} - (bt)^{p_2} e^{-bt}, \quad (3.32)$$

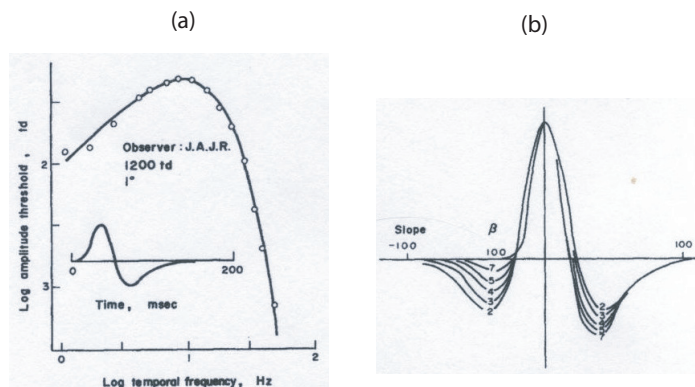
wobei der Index w in h_w andeuten soll, dass es sich um Watsons Alternative handelt. Watson schätzte die freien Parameter a , b , p_1 und p_2 als $a = 1/4.94$, $b = 1/6.58$, $p_1 = 8$ und $p_2 = 9$.

In einem zweiten Schritt zeigt Watson nun, dass diese Impulsantwort resultiert, wenn man (i) vom "richtigen" Entdeckensmodell ausgeht, dass auf der Annahme zeitlicher Wahrscheinlichkeitssummation beruht, und (ii) dazu die von ihm vorgeschlagene Formel (??) benutzt; sie sei hier zur leichteren Bezugnahme noch einmal angegeben:

$$\psi_w(c) = 1 - \exp \left[- \int_{-\infty}^{\infty} |\alpha g(t; c)|^\beta dt \right]. \quad (3.33)$$

⁴Für Ingenieure ist der Unterschied zwischen " $\approx$ " und " $=$ " anscheinend oft vernachlässigbar, vor allem, wenn der Erfolg ihnen mit dieser Gleichsetzung Recht gibt.

Abbildung 3.10: Watsons (1981) Alternativen – (a) Mit Amplitudenspektrum kompatible Impulsantwort, (b) Erklärung der Roufs und dBlommaertschen Impulsantwort durch zeitliche Wahrscheinlichkeitssummation



Die psychometrische Funktion ist wieder mit w indiziert worden, um sie als Watsons Funktion zu kennzeichnen. Hier taucht noch der freie Parameter α auf, der die Position bzw. den Bereich von ψ_w auf der c -Skala bestimmt, der aber in die Funktion g absorbiert werden kann und deshalb im Folgenden weggelassen wird. Zunächst kann festgehalten werden, dass der Integrationsbereich auf $[0, \infty)$ beschränkt werden kann, denn g ist für $t < 0$ gleich Null. Der für Watsons Argument wichtige Parameter ist β . Der Wert von β bestimmt die Steilheit der Funktion: je größer β , desto steiler ist ψ_w . Nun sind psychometrische Funktionen des Typs (3.33) in der Folge von Quicks (1974) Arbeit vielfach auf spatiale wie auch auf zeitliche Fragestellungen angewendet worden, wobei für β Werte in der Größenordnung zwischen 2 und 8 geschätzt wurden; Watson (1979) gibt an, dass von 104 β -Schätzungen 89 zwischen 3 und 7 liegen. Roufs und Blommaert (1981) schätzten für ihre psychometrischen Funktionen (die aber nicht zum Zwecke der Interpretation für zeitliche Wahrscheinlichkeitssummation geschätzt wurden) einen Wert von $\hat{\beta} = 6.6$, den Watson aber für eine Fehlschätzung hält, wobei er auf eine Arbeit von Nachmias (1981) verweist, in der die Probleme der Schätzung von β diskutiert werden.

Watson (1981) reproziert nun die Schätzungen der Impulsantwort von Roufs und Blommaert, indem er einerseits die Impulsantwort (3.32) zugrundelegt, und andererseits annimmt, dass Entdecken dem W-Summationsprinzip⁵ folgt, und dazu die Perturbationsmethode simuliert. Für verschiedene Werte von β ergeben sich nun verschiedene Schätzungen (Reproduktionen) der Impulsantwort. Die Resultate werden in Abb. 3.10, (b) gezeigt (Watson (1981)). Für die "vernünftigen" Schätzungen $\hat{\beta} = 2$ bis $\hat{\beta} = 7$ ergeben sich stets W-Summationseffekte, die sich als inhibitorische Dellen im ersten Teil der mit der Perturbationsmethode *geschätzten* Impulsantwort zeigen, obwohl die "wahre", also im Simulationsprogramm zugrundegelegte Impulsantwort diese Delle nicht aufweist. Erst für den Wert $\beta = 100$ verschwindet die Delle

⁵W-Summation = Wahrscheinlichkeitssummation

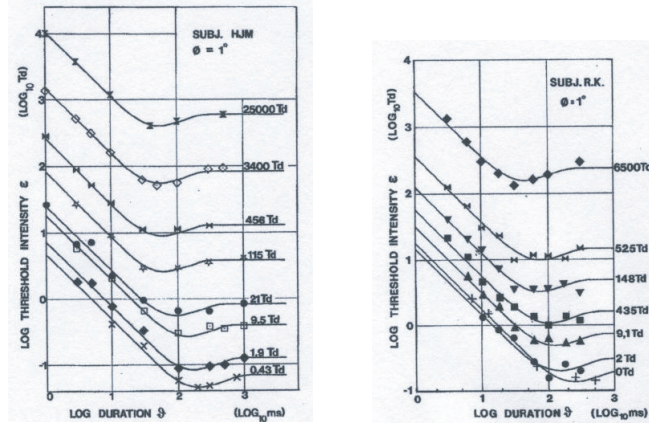
und die "wahre" Impulsantwort wird reproduziert. Damit ist, so Watson (1981) , gezeigt worden, dass die von Roufs und Blommaert vorgeschlagene Impulsantwort für phasische Kanäle, wie sie in Abb. 3.8 (a) gezeigt wird, ein durch Nichtbeachtung der zeitlichen W-Summation erzeugtes Artefakt ist.

Blommaert und Roufs (1987) haben ihr Modell verteidigt. Dazu legen sie Eine Reihe von neuen Daten vor, die die Ergebnisse der ersten Arbeit von 1981 bestätigen und erweitern. Die von Watson (1979) und (1981) behauptete psychometrische Funktion wird aber nicht in Frage gestellt, Blommaert und Roufs akzeptieren sie als eine Möglichkeit, über die Schätzung des Parameters β den Effekt der W-Summation abzuschätzen. Denn dieser Parameter bestimmt die Varianz der zufälligen Veränderungen, deren Verteilung die Weibull-Verteilung ist: je größer der Wert von β , desto kleiner ist diese Varianz. Blommaert und Roufs merken an, dass die psychometrische Funktion um so flacher ist – das heißt, dass sie einen um so kleineren β -Wert hat –, je länger es dauert, die psychometrische Funktion zu schätzen. Der kleinere β -Wert sei aber ein Artefakt der Methode, mit der die psychometrische Funktion bestimmt wird. Der Weibull-Ansatz, auf die Kritik Watsons basiert, wird dabei nicht in Frage gestellt, sondern nur das Argument, dass der β -Wert, den

1. Je länger Zeitintervall, desto flacher die psychometrische Funktion (his section 5.2 on the nature of the noise)
2. The lower frequency components of the noise play a minor role; the only information about the noise is in the psychometric function.
This is an intuitive argument, there is no exact development of the argument. Roufs bezieht sich ebenfalls auf die Weibull-Funktion, β ist der "noise parameter".
3. Es gäbe eine Überschätzung der W-Summationseffekte, wenn man nur auf die Psychom Funktion fokussiere. Die Überschätzung sei um so größer, je länger es daure, die Funktionen zu schäten.
4. Also argumentieren B + R, dass ein großer β -Wert der richtige ist: man müsse "kurz" messen, weil dann der Rauscheffekt klein sei.

Neben der Steigung der psychometrischen Funktion: die Impulsantwort kann benützt werden, die Schwellenintensitäten für zeitliche Rechteckstimuli verschiedener Dauer vorherzusagen. Für Dauern bis zu einer bestimmten Dauer T_b gilt das Blochsche Gesetz $c_p T_s = k$ eine Konstante, wobei c_p die Schwellenintensität und $T_s < T_b$ die Dauer der Präsentation ist. Nun ist das Blochsche Gesetz schon in Abb. 3.3 getestet worden: für $T_s > T_b$ sinkt c_p nicht mehr. Aber dieser Befund hängt vom Stimulus ab: für kleine Kreisscheiben gilt der Befund von Abb. 3.3, – es wird ein tonischer Kanal angesprochen. Für größere Kreisscheiben fand Roufs bereits 1974 einen anderen: es zeigte sich in den Plots $\log T_s$ versus $\log c_p$ eine Delle ("dip"). Sagt man nun mit der Impulsantwort für phasische Kanäle diesen Plot vorher, so wird genau dieser "dip" vorhergesagt, wobei die Annahme des Peak-Detection eingeht. Blommaert und Roufs (1987) haben diese Experimente für verschiedene Hintergrundhelligkeiten wiederholt. Falle Versuchspersonen ergibt sich das gleiche Bild. Abb. 3.11 zeigt die

Abbildung 3.11: Schwellenintensitäten für Rechteckstimuli, transients Kanal. Aus. Blommaert und Roufs (1987)



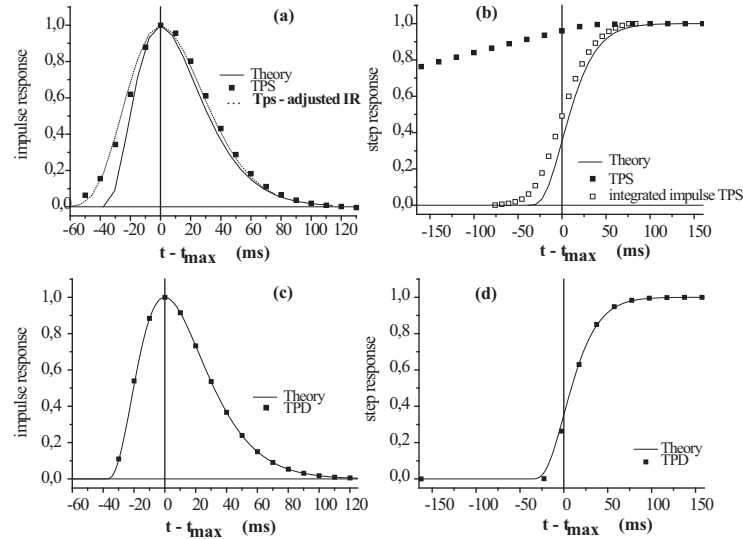
Daten für zwei Versuchspersonen, wobei die durchgezogenen Linien die Vorhersagen anhand der Impulsantworten unter Annahme des Peak-Detection-Mechanismus sind. Diese Vorhersagen stützen sowohl die Hypothese der Linearität im Schwellenbereich wie auch die des Peak-Detection.

Natürlich kann man fragen, ob es nicht doch möglich sei, dass diese Ergebnisse auch durch ein Modell der zeitlichen W-Summation vorhergesagt werden können. Man sollte hier nicht an das Watsonsche Modell denken, da die Diskussion dieses Modells gezeigt hat, wie problematisch es ist. Es sollte ein Modell sein, dass zumindest realistische Annahmen über die Autokorrelation des Rauschens macht. Dies ist das Modell (2.89), d.h.

$$\psi(c) = 1 - \exp \left[-\frac{\sqrt{\lambda_2}}{2\pi} \int_0^T \exp \left(-\frac{1}{s} (S - g(t; c))^2 \right) dt \right]. \quad (*)$$

Wie bereits angemerkt, hängt die Geschwindigkeit, mit der das Rauschen fluktuiert und die damit die Autokorrelation bestimmt, vom Wert des freien Parameters λ_2 ab. Eine Möglichkeit, die Hypothese des Entdeckens durch zeitliche W-Summation zu testen, ist, die Perturbationsmethode mit fälschlich angenommenem Peak-Detection anzuwenden. Die Schwellen werden also bestimmt, indem c_p über (2.89), d.h. also (*) bestimmt werden, dieser c_p -Wert aber so verrechnet wird, als gelte Peak-Detection. Die zu bestimmenden Antwortfunktionen seien die Impuls- und die Schrittwort für einen tonischen Kanal, wie sie in Abb. 3.9 gezeigt werden. Sollte das Entdecken dem zeitlichen W-Summationsprinzip entsprechen, so müsste sich dies insbesondere bei der Bestimmung der Schrittwort zeigen, die ja von einem bestimmten Zeitpunkt an konstant einen Wert an der Schwelle S hat. Zeitliche W-Summation sollte mit größer werdender Stimulusdauer den Wert von c_p reduzieren. Ein Test dieser Art wurde von Mortensen (2007b) vorgestellt. Abb. 3.12 zeigt das Ergebnis. (a) zeigt die Impulsantwort, wobei die Quadrate die Schätzungen zeigen, wenn in Wahrheit das

Abbildung 3.12: Test des Modells der zeitlichen W-Summation. Aus Mortensen (2007b)



Entdecken durch zeitliche W-Summation geschieht. Die Impulsantwort wird zumindest qualitativ korrekt dargestellt, allerdings ist sie breiter als die tatsächliche Impulsantwort. (b) zeigt die Schrittantwort. Die Quadrate sind die Vorhersagen unter der Bedingung der zeitlichen W-Summation. Diese Punkte hätte die Perturbationsmethode liefern müssen, wären die Versuchspersonen einer W-Summationsstrategie gefolgt. (c) und (d) zeigen die Impuls- und Schrittantwort, wie sie durch die Perturbationsmethode erzeugt werden, wenn das Entdecken tatsächlich Peak-Detection ist. Die "wahren" Funktionen werden perfekt reproduziert.

Das Interessante an diesem Test ist, dass der Wert von λ_2 offenbar keine Rolle spielt. Alle für verschiedene λ_2 -Werte berechneten Vorhersagen sind identisch. Dies bedeutet, dass die Perturbationsmethode von der Autokorrelation unabhängige Schätzungen der zu erfassenden Funktion liefert, falls das Entdecken tatsächlich durch zeitliche W-Summation geschieht. Die einzige Einschränkung der Allgemeinheit, die für das Modell (2.89), d.h. also (*), gilt, ist, dass das Rauschen stationär ist; diese Annahme ist einerseits üblich, andererseits liegt kein Modell mit nicht-stationärem Rauschen vor, bei dem $\psi(c)$ in geschlossener Form angegeben werden kann. Dies bedeutet nicht, dass man sich mit dem Modell (*) zufriedengeben muß, allerdings zeigen relativ allgemeine Berechnungen zur Aktivität von Neuronenpopulationen, dass zeitliche W-Summation keineswegs als allgemeines, kanonisches Modell angenommen werden müßte (Mortensen, 2007b). Es scheint eher so zu sein, dass die tatsächlich wirkenden Entdeckensmechanismen sich an der durchschnittlichen Aktivität, wie sie durch die Systemantwort $g(t)$ repräsentiert wird, orientieren. Insbesondere scheint es so zu sein, dass Versuchspersonen in der Lage sind, auf den Maximalwert von g zu fokussieren, und dass dabei die Dauer, die g diesen Wert

annimmt (wie bei der Schrittantwort) irrelevant ist.

3.4 Systemidentifikation im Ortsbereich

3.4.1 Das visuelle System als einzelner Kanal

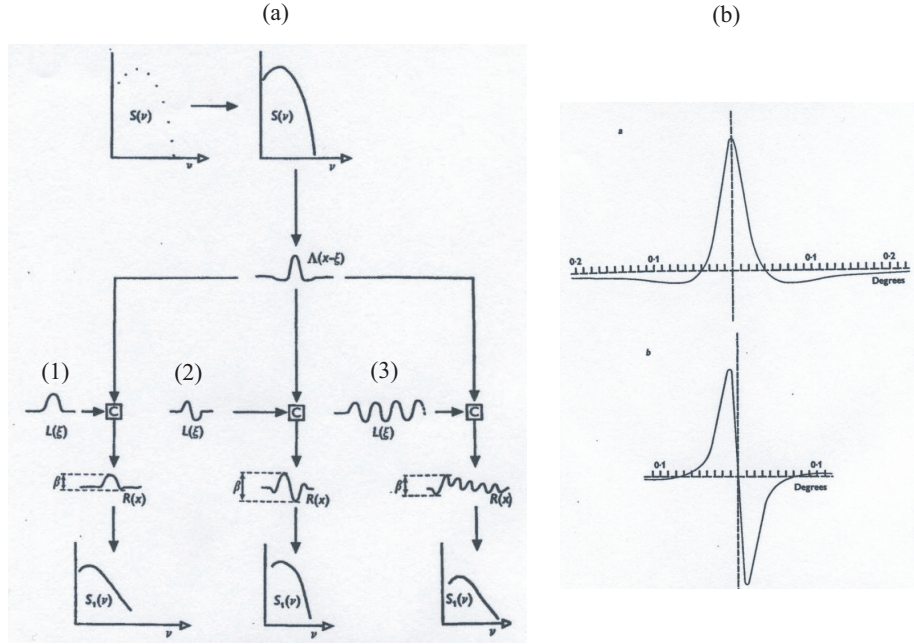
Erstmal Versuch eines Überblicks:

1. Transferfunktionen im Ortsbereich: Kelly-Arbeiten, Campbell & Kulikowski etc (1966); Campbell & Robson – Square wave gratings and Fourieranalysis
2. Erklärung von "Gesetzen" – Ricco?
3. Bestimmung von Impulsantworten
4. Fourier-Theorien: Marcelja, Fiorentini, Adaptationsdaten, Sachs et al, Graham 1977,
5. Probleme mit dem Fourier-Ansatz: rezeptive Felder endlich etc – Gabor-Muster
6. ein Paar spezielle Sachen: Hartmann (rekursive Strukturen), Hans du Buf und Machbänder
7. Features, Matched Filter, Lernen und Selbstorganisation; includes Kulikowski, Hauske etc
8. Texturen
9. Neuronale Netze als dynamische Systeme? Gehört vielleicht noch in den Systemtheorieteil.

Die Systemtheorie eignet sich natürlich auch dazu, die Übertragungseigenschaften im Ortsbereich zu charakterisieren. Insbesondere kann man untersuchen, ob das visuelle System über spezifische "Kanäle" oder Filter – also neuronale Teilsysteme – zur Identifikation spezifischer Musterelemente verfügt, deren Aktivierung insgesamt dann ein Gesamtperzept des Musters ergibt. Sofern diese Kanäle im Schwellenbereich linear operieren, kann man dann versuchen, diese Kanäle durch ihre jeweiligen Transfer- oder Systemfunktionen bzw. die Impulsantworten im Ortsbereich zu bestimmen, was dazu führt, dass Sinusgitter wichtige Stimulismuster darstellen.

Die Anzahl der Arbeiten zu diesem Thema ist zu groß, als hier einen vollständigen Überblick geben zu können. Es sollen nur die wichtigsten empirischen Ergebnisse und Modelle vorgestellt werden. Westheimer (1960) war sicher nicht der erste, der die Empfindlichkeit für Sinusgitter ausmaß, ohne allerdings weitere Spekulationen über neuronale Mechanismen anzustellen. Campbell und Gubisch (1966) vermessen die optische Qualität des menschlichen Auges und diskutieren die Hypothese, dass die neuronalen Verbindungen in der Retina die vermutete schlechte Qualität der Linse kompensieren sollen: sie befinden, dass die Linse gar nicht so schlecht sei und der

Abbildung 3.13: Campbell, Carpenter und Levinson (1969):(a) Stimuli und Vorher-
sagen, (b) Impulsantwort und Derivierte der Impulsantwort



Korrekturmechanismus entweder nicht existiert oder, wenn doch, nur eine geringe Rolle spielt.

Campbell, Carpenter und Levinson (1969) haben die mathematischen Implikationen der Hypothese, dass das visuelle System als in erster Näherung lineares System beschrieben werden kann, exploriert. So bestimmen sie die Impulsantwort aus geschätzten Transferfunktionen und bestimmten daraus die Reaktion auf aperiodische Muster. Sie schlagen

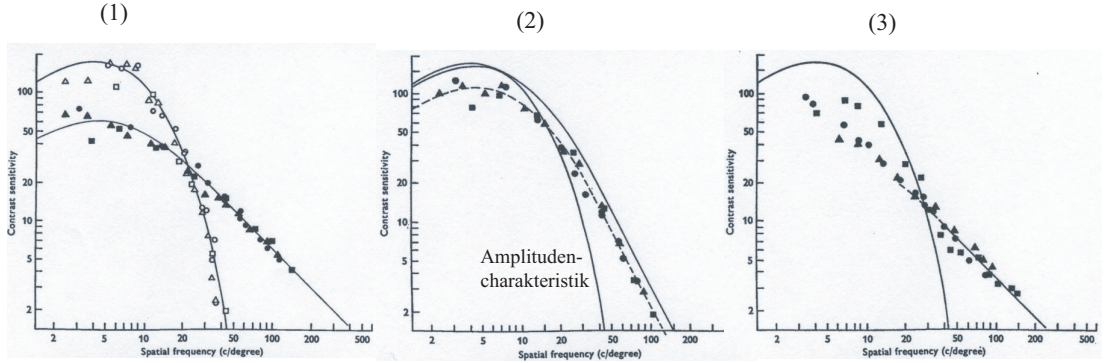
$$h(x) = \frac{K}{\pi} \left(\frac{\alpha}{\alpha^2 + x^2} - \frac{\beta}{\beta^2 + x^2} \right) \quad (3.34)$$

vor. Damit kann man die Antwort auf einen beliebigen Stimulus, der durch die Luminanzverteilung L definiert ist, über das Faltungsintegral

$$g(x) = \int_{-\infty}^{\infty} h(x - \xi) L(\xi) d\xi. \quad (3.35)$$

betimmt werden. Campbell et al. betrachten drei Stimuli: einen einzelnen vertikalen Streifen ("bar"), definiert durch einen halben Zyklus einer Sinusfunktion ((1) in Abb. 3.13), einen "Doppelstreifen" ("double bar"), der durch einen vollen Zyklus einer Sinusfunktion definiert ist ((2) in Abb. 3.13), und die Kante eines Sinusgitters (von der Mitte des Bildschirms bis zum Rand des Feldes, (3) in Abb. 3.13). Stimulus (1)

Abbildung 3.14: Campbell, Carpenter und Levinson (1969): Daten für die Stimuli (1), (2) und (3)



insbesondere durch

$$L(\xi) = \begin{cases} \cos \Omega \xi, & -1/4\nu \leq \xi \leq 1/4\nu \\ 0, & \text{sonst} \end{cases} \quad (3.36)$$

Der Stimulus ist eine gerade Funktion, und damit ist die Antwort auf ihn durch

$$g(x) = a \int_{-1/4\nu}^{1/4\nu} \nu h(x - \xi) \cos \Omega \xi d\xi. \quad (3.37)$$

Nach partieller Integration liefert die Taylor-Entwicklung von g

$$g(x) \approx 2a \left[\frac{h(x)}{\Omega} - \frac{h''(x)}{\Omega^2} + \dots \right], \quad h'' = dh^2/dx^2 \quad (3.38)$$

Für hinreichend großen Wert von ν sind die höheren Ableitungen vernachlässigbar. a ist der Schwellenkontrast, der für größere Werte von ν proportional zu ν wird, – womit Riccos Gesetz für schmale Streifen vorhergesagt wird. Campbell et al nehmen zwei Möglichkeiten für den Entdeckensmechanismus an: einfache spatiale Peak-Detektion – hier wird der Stimulus entdeckt, wenn $g_{\max} = \max_x g(x)$ einen kritischen Wert erreicht, oder Peak-toPeak-Detektion – dann muß $\Delta g = \max_x g(x) - \min_x g(x)$ einen kritischen Wert erreichen. Abb. 3.14 zeigt die Daten für drei Versuchspersonen. Die hellen Symbole repräsentieren die System- bzw. die Amplitudenfunktion des visuellen Systems, die Daten die Voraussagen anhand von Stimulus (1). Die durchgezogenen Kurven entsprechen der Theorie, die sich hier also als mit den Daten kompatibel erweist. Für Stimulus (2) ergeben sich Schwierigkeiten, weil nicht klar ist, ob sich die Vpn nach der $\max_x g(x)$ oder nach der Δg -Regel verhalten. Die äußere durchgezogene Linie entspricht den Vorhersagen aufgrund der Δg -Regel. Wählen die Vpn mal die $g_{\max}$ – und mal die Δg -Regel, so ergibt sich im Durchschnitt die gestrichelte Vorhersage, die mit den Daten der drei Vpn sehr gut übereinstimmt. Auch für

Stimulus (3) entsprechen die Daten den Vorhersagen. Insgesamt haben die Autoren demonstriert, dass die Linearitätsannahme mit einem Peak-Detection-Mechanismus ein durchaus vernünftiger Ansatz ist, die Struktur des visuellen Systems zu erfassen.

Maffei und Fiorentini (1973) haben gezeigt, dass insbesondere die *simple cells* des visuellen Kortex die Eigenschaft von Ortsfrequenzfiltern haben.

3.4.2 Die direkte Schätzung der Impulsantwort im Ortsbereich

Wenn das visuelle System als ein einziger Kanal konzipiert werden kann, liegt es nahe, diesen Kanal durch seine spatio-temporale Impulsantwort zu charakterisieren; kennt man diese, so läßt sich die Reaktion des Kanals – also des visuellen Systems – für einen beliebigen Stimulus vorhersagen, zumindest wenn der Kontrast im Schwellenbereich bleibt. Die Frage ist, ob die Impulsantwort für jede retinale Position identisch ist.

Der Einfachheit halber werde zunächst angenommen, dass die Retina ein radial-symmetrisches System ist. Für die (spatiale) Impulsantwort $h_s(x, y)$ gilt dann

$$h_s(x, y) = h_s(\sqrt{x^2 + y^2}). \quad (3.39)$$

Es genügt dann, h nur bezüglich der x -Achse zu bestimmen; man erhält dann die LSF (*Line Spread Function*⁶)

$$LSF(x) = h(x) = \int_{-\infty}^{\infty} h_s(x, y) dy. \quad (3.40)$$

Die LSF ist die Antwort auf eine sehr schmale senkrechte Linie, idealerweise eine Dirac-Delta-Funktion $\delta(x)$. Tatsächlich wird man als Luminanzprofil nur ein Rechteck- oder Gaußprofil erzeugen können. Für das Rechteckprofil ergibt sich dann für die Position $x = \pm\varepsilon$ die Approximation

$$s(x, \varepsilon) = \begin{cases} \Delta x \delta(x \pm \varepsilon), & \pm\varepsilon - \Delta x/2 \leq x \leq \pm\varepsilon + \Delta x/2 \\ 0, & \text{sonst} \end{cases} \quad (3.41)$$

Die LSF an der Position ε läßt sich relativ leicht bestimmen, wenn (i) eine vertikale Linie in ε mit dem Kontrast c präsentiert und simultan im Abstand $\pm a$ zwei weitere, flankierende Linien mit dem Kontrast c/s , und wenn man weiter davon ausgeht, dass nach Maßgabe des Prinzips der spatialen sowie der zeitlichen Peak-Detection entdeckt wird. Um mögliche Verfälschungen durch zeitliche transitorische Effekte zu vermeiden, sollte der Stimulus "sanft" ein und wieder ausgeschaltet werden. Die zeitliche Präsentation kann etwa durch eine Gaußfunktion der Form

$$s_t(t) = \exp \left[-\frac{(t - \mu_t)^2}{2\sigma_t^2} \right] \quad (3.42)$$

⁶Die deutschen Ausdrücke sind hier ein wenig umständlich: h_s heißt *Punktverwaschungsfunktion*, was dem englischen *point spread function* entspricht. Da sich mittlerweile die englischen Abkürzungen durchgesetzt haben, wird hier ebenfalls von ihnen Gebrauch gemacht.

definiert werden. Insgesamt ist dann das Stimulismuster durch

$$L(x, t) = L_m(1 + cs_\varepsilon(x, a)s_t(t)) \quad (3.43)$$

definiert, wobei $s_\varepsilon(x, a)$ aus einer vertikalen in Linie $x = \varepsilon$ und zwei flankierenden vertikalen Linien an der Position $\varepsilon + a$ und $\varepsilon - a$, die nur den halben Kontrast wie die zentrale Linien haben. $L(x, t)$ ist die Luminanz an der Stelle x zur Zeit t , L_m ist die mittlere Luminanz, und

$$c = \frac{L_{\max} - L_{\min}}{L_m}. \quad (3.44)$$

Läßt man der Übersichtlichkeit wegen die Abhängigkeit von der Zeit weg, so ergibt sich die Antwort $g(x)$ speziell für $\varepsilon = 0$, also im Zentrum der Fovea, durch die Faltungen

$$\begin{aligned} g(x) = & c_a \Delta x \int_{-\infty}^{\infty} h(x - \xi) \delta(\xi) d\xi \\ & + \frac{c_a \Delta x}{2} \left[\int_{-\infty}^{\infty} h(x - \xi) \delta(-a - \xi) d\xi + \int_{-\infty}^{\infty} h(x - \xi) \delta(a - \xi) d\xi \right] \end{aligned} \quad (3.45)$$

wo jetzt c_a statt c geschrieben wurde, um anzudeuten, dass der Kontrast für den Flankenabstand a gemeint ist. Um die Ausdrücke zu vereinfachen, wird nun der Faktor Δx in die Definition von h absorbiert, denn h kann experimentell sowieso nur bis auf einen Faktor bestimmt werden. Es folgt dann

$$g(x) = c_a h(x) + \frac{c_a}{2} [h(x + a) + h(x - a)] \quad (3.46)$$

folgt. Da die mittlere Linie in $\varepsilon = 0$ präsentiert wird und man also die LSF für das Zentrum der Fovea bestimmt, kann angenommen werden, dass die Antwort auf das Linienmuster hier ihr Maximum hat, also $h(0) = \max_x h(x)$ und dementsprechend $\max_x g(x) = g(0)$. Die postulierte Radialsymmetrie von h_x impliziert, dass h symmetrisch in Bezug auf $\varepsilon = 0$ verläuft, so dass $h(-a) = h(a)$ ist. Weiter impliziert die Peak-Detection-Annahme, dass die mittlere Linie entdeckt wird, wenn $g(0) = k$, k eine bestimmte Konstante, ist; die flankierenden Linien werden nicht entdeckt, da sie ja nur den halben Schwellenkontrast der mittleren Linie haben. Also folgt für den Schwellenkontrast c_a

$$g(0) = k = c_a h(0) + \frac{c_a}{2} [h(a) + h(-a)], \quad (3.47)$$

und wegen $h(-a) = h(a)$ hat man

$$k = c_a [h(0) + h(a)]. \quad (3.48)$$

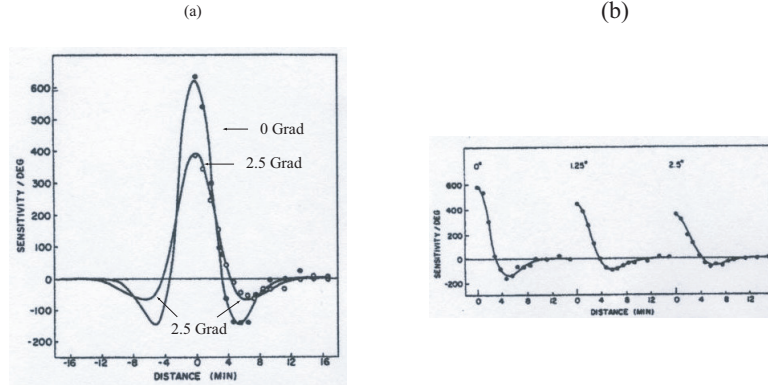
Präsentiert man nun die mittlere Linie *ohne* flankierende Linien, so werde sie entdeckt, wenn ihr Kontrast gleich c_0 ist:

$$k = c_0 h(0), \quad (3.49)$$

woraus

$$S_0 = \frac{1}{c_0} = \frac{1}{k} h(0) \quad (3.50)$$

Abbildung 3.15: Line Spread Funktionen (Hines, 1976) für verschiedene retinale Positionen. In (a) sind die LSFen für die Position 0° (Fovea) übereinandergezeichnet; die LSF für 0° hat einen größeren Maximal- und einen kleineren Minimalwert als die LSF für 2.5° Exzentrizität. (b) zeigt noch einmal die LSFen für 0° , 1.25° und 2.5° . Die durchgezogenen Funktionen sind jeweils Differenzen zweier Gauss-Funktionen.



folgt; S_0 ist die *Sensitivität* für die mittlere Linie, wenn sie allein präsentiert wird, und $S(a) = 1/c_a$ ist dementsprechend die Sensitivität für die Linie plus die flankierenden Linien an den Stellen $\pm a$. Die Gleichung (3.48) impliziert dann

$$\frac{1}{c_a} - \frac{1}{c_0} = \frac{1}{k} h(a). \quad (3.51)$$

Anmerkung: Es ist

$$c_0^* = \lim_{a \rightarrow 0} c_a = \frac{1}{2} c_0. \quad (3.52)$$

Für hinreichend kleinen Wert von a lässt sich $h(a)$ in der Form

$$h(a) = h(0) + ah'(0) + o(a)$$

schreiben, wobei $o(a)$ eine Funktion von a ist, die schneller als a gegen Null geht. Also erhält man aus (3.48)

$$k = c_a[h(0) + h(0) + o(a)] = c_a[2h(0) + o(a)],$$

und

$$k = \lim_{a \rightarrow 0} c_a[h(0) + h(0) + o(a)] = c_0^* 2h(0) = c_0^* \frac{2k}{c_0} \quad (3.53)$$

oder

$$c_0 = 2c_0^*,$$

entsprechend (3.52).

Die Bestimmung des LSFen h_ε an den Positionen $\varepsilon \neq 0$ ist analog. Abb. 3.15 zeigt LSFen für verschiedene Exzentrizitäten. Offenbar werden sie um so flacher, je größer

die Exzentrizität ist. Die Meßpunkte (vergl. (??)) entsprechen sehr gut einer als Differenz zweier Gauss-Funktionen definierten Funktion:

$$h_\varepsilon(x) = \frac{A}{\alpha\sqrt{2\pi}} \exp\left[-\frac{(x-\varepsilon)^2}{2\alpha^2}\right] - \frac{B}{\beta\sqrt{2\pi}} \exp\left[-\frac{(x-\varepsilon)^2}{2\beta^2}\right]. \quad (3.54)$$

Die Werte der freien Parameter A und B sind spezifisch für die einzelnen Versuchspersonen, das Verhältnis β/α erwies sich allerdings für alle Vpn und Exzentrizitäten als konstant: $\beta/\alpha = 1.5$.

Läßt sich nun das visuelle System tatsächlich als ein 1-Kanalsystem konzipieren und gilt die Hypothese der Linearisierbarkeit für Kontraste im Bereich der Schwelle, so sollten sich die Kontrastschwellen für nachgerade beliebige Stimuli durch die Faltung des korrespondierenden Luminanzprofils mit der lokalen LSF (3.54) vorhergesagen lassen. Hines wählte verschiedene "Balkenmuster" aus:

$$s(x) = \begin{cases} \cos(\pi(x-\varepsilon)/W), & |x-\varepsilon| < W/2, & \text{cosine bar} \\ 1 - \cos(\pi(x-\varepsilon)/W), & 0 < x-\varepsilon < W, & \text{cosine edge} \\ 1, & x \in [\varepsilon \pm d/2], & \text{square bar} \end{cases} \quad (3.55)$$

Die Ergebnisse der Messungen und Vorhersagen werden in Abb. 3.16 gezeigt. Die durchgezogenen Kurven sind mit der LSF (3.54) berechnet worden. Offenbar sind für diese Muster die Vorhersagen sehr befriedigend.

Die Ergebnisse für square bars sind ernüchternd. Der mangelnde Fit für diese Reizmuster wird von Hines erklärt durch die Tatsache, dass bei hinreichend breiten Mustern die Verschiedenheit der LSFen für jeden x -Wert berücksichtigt werden müsse. Eine andere Möglichkeit ist, "size-tuned" Mechanismen anzunehmen, wie sie von Thomas (1979) vorgeschlagen wurden. Eine weitere Alternative sind Matched Filter, die in Abschnitt 3.4.6 behandelt werden. Eine weitere Möglichkeit ist, dass bei einem Rechteckmuster die Kanten unabhängig voneinander entdeckt werden und Wahrscheinlichkeitssummationseffekte beim Entdecken dieser Muster auftreten (Kulikowski & Kong-Smith, 1975).

Statt einzelner cosine bars kann man natürlich auf Kosinusgitter präsentieren, wobei die Anzahl der Zyklen eine für das Entdecken des Gitters wichtige Variable wird. Abb. 3.17, (b) zeigt die Daten. Die durchgezogenen Kurven sind die Vorhersagen des 1-Kanalmodells anhand der LSF (3.54). Die Abhängigkeit von der Anzahl der Zyklen wird offenbar nicht korrekts vorhergesagt. Alternative Modelle, wie sie im Zusammenhang mit dem Rechteckbalken bereits genannt wurden, müssen zur Erklärung der Daten herangezogen werden: ein einfaches 1-Kanalmodell ist offenbar kein hinreichendes Modell für alle Daten.

3.4.3 Neuronale Kanäle: Orientierungen

Campbell, Kulikowski und Levinson (1966) zeigen Gratings verschiedener Ortsfrequenz in verschiedenen Orientierungen um die bereits bekannte Abhängigkeit der

Abbildung 3.16: Messungen und ihre Vorhersagen: cosine bar und cosine edge für $\varepsilon = 0^0$ und $\varepsilon = 2.5^0$ (Position der Unstetigkeit). Auf der x -Achse sind die Breiten der Balken angegeben. Gefüllte Kreise: 0^0 , offene Kreise: 2.5^0 .

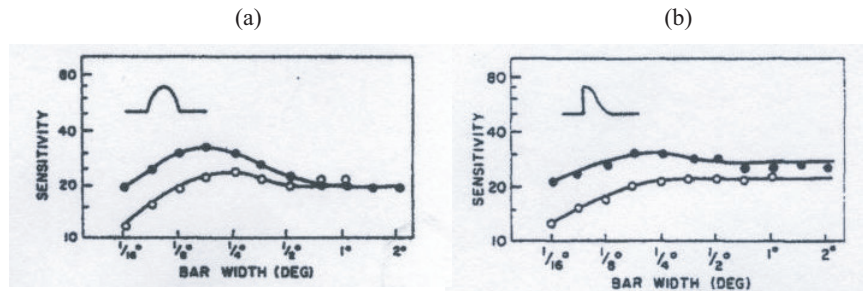
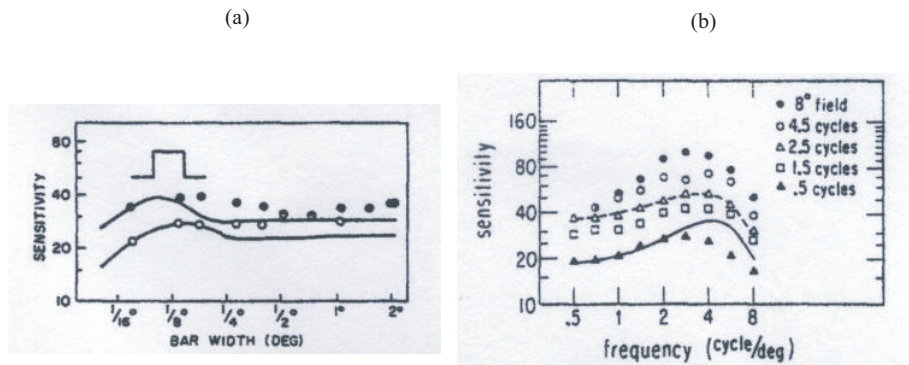


Abbildung 3.17: Messungen und ihre Vorhersagen: square bars an den Stellen 0^0 und 2.5^0 . Auf der x -Achse sind wieder die Breiten der Balken angegeben. Gefüllte Kreise: 0^0 , offene Kreise: 2.5^0 .



Wahrnehmung von der Orientierung zu untersuchen und kommen zu dem Schluß, dass diese Abhängigkeit nicht durch Eigenschaften der Optik, sondern der neuronalen Strukturen des visuellen Systems bedingt ist. Hubel und Wiesel (1959, 1962) haben durch ihre grundlegenden Arbeiten gezeigt, dass im primären visuellen Kortex Zellen existieren, die auf Linien mit bestimmten Orientierungen reagieren. Es liegt nahe, zu vermuten, dass diese Zellen eine erste Kodierung von visuellen Mustern vornehmen. Man kann Gittermuster als Muster, die aus parallelen Linien bestehen, auffassen, vor allem, wenn sie aus weißen und schwarzen "Balken" bestehen. Aber auch Sinusgitter kann man Linienmuster interpretieren. In diesem Sinne haben Campbell und Kulikowski (C & K) (1966) Sinusgitter interpretiert, um die Orientierungselektivität des visuellen Systems psychophysisch zu untersuchen. Zu diesem Zweck führten sie ein Superpositionsexperiment durch. C & K erzeugten jeweils ein Sinusgitter auf zwei Oszilloskopen; die Ortsfrequenz der Gitter war jedesmal 10 c/deg (Zyklen pro Grad Schwinkel, cycles per degree). Der Kontrast der Gitter wird

Abbildung 3.18: Kontrastdefinition für ein Sinusgitter

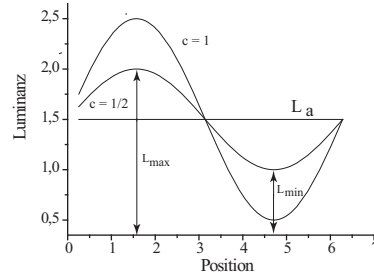
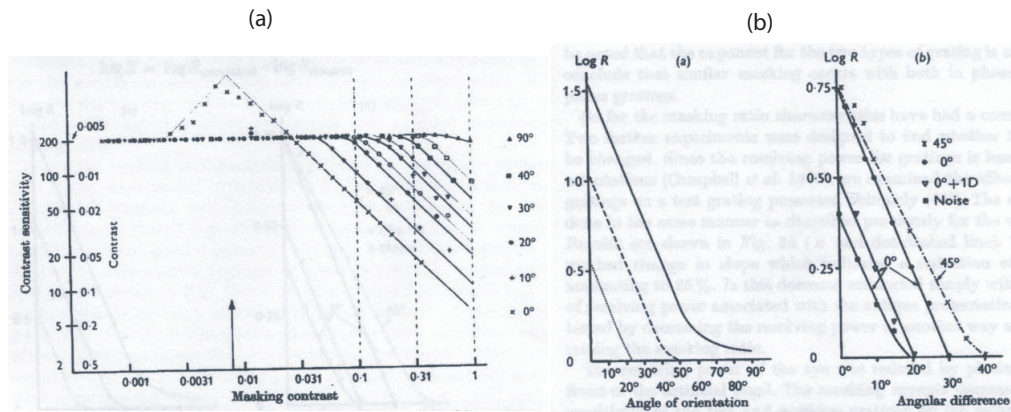


Abbildung 3.19: Orientierungskontraste (Campbell & Kulikowski (1966))



entsprechend der Definition

$$c = \frac{L_{\max} - L_{\min}}{L_{\max} + L_{\min}} \quad (3.56)$$

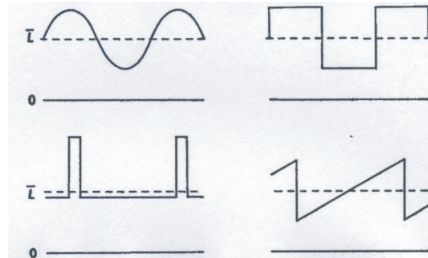
berechnet (vergl. Abb. 3.18). Eines der Gitter ist das Testgitter, das andere das Maskierungsgitter. Der Kontrast des Testgitters konnte manipuliert werden, der des Maskierungsgitters war für eine Versuchsbedingung konstant; das Maskierungsgitter wurde mit fixer Orientierung optisch superponiert. Die Orientierung des Testgitters wurde nun variiert und für jede Orientierung (0° , 10° , 20° , 30° , 40° , 90°) wurde die *Kontrastsensitivität* $S = 1/c_0$ bestimmt, wobei c_0 die Kontrastschwelle für das Testgitter ist. C & K berechneten dann einen Maskierungsquotienten

$$R = \frac{S_{\text{unmaskiert}}}{S_{\text{maskiert}}} \quad (3.57)$$

Abb. 3.19 zeigt die Ergebnisse; in (b) sind die Ergebnisse schematisch, aber dafür deutlich dargestellt: es wird

$$\log R = \log_{10} S_{\text{unmaskiert}} - \log_{10} S_{\text{maskiert}}$$

Abbildung 3.20: Luminanzprofile der Gitter in der Untersuchung von Campbell & Robson (1968))



gegen den Orientierungswinkel aufgetragen. Die durchgezogene, die gestrichelte und die punktierte Kurve entsprechen verschiedenen Maskierungsniveaus (Kontraste 1, .31 und .1). Die Daten legen nahe, dass die Linien"kanäle" (später werden diese Kanäle als Gitterkanäle aufgefat) orientierungsempfindlich sind. G & K diskutieren ihre Ergebnisse in Bezug auf die Befunde Hubel und Wiesel (1965): Stimuli, die eine 90° -Orientierung in Bezug auf die optimale Orientierung einer Zelle haben, beeinflussen die Aktivität dieser Zelle nicht mehr; nur Orientierungsabweichungen bis zu maximal 30° können die Aktivität der Zelle beeinflussen. G & K finden eine bemerkenswerte Übereinstimmung von psychophysischen Resultaten und physiologische Resultaten, die am visuellen Kortex der Katze gewonnen wurden.

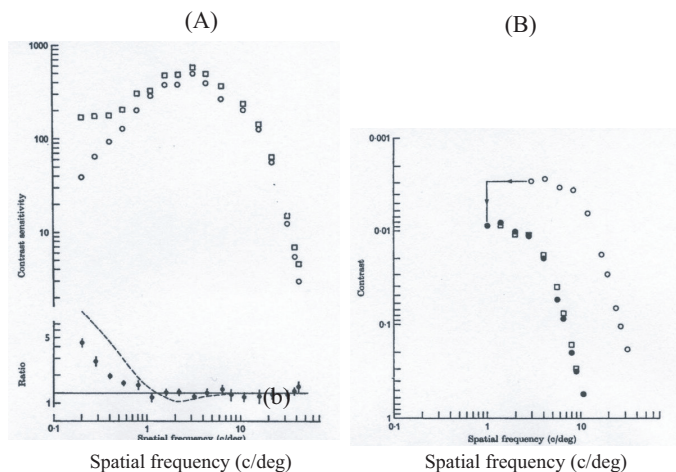
3.4.4 Neuronale Kanäle: Ortsfrequenzen

Campbell und Robson (1968) präsentierten Gitter, deren Luminanzprofile in Abb. 3.20 gezeigt werden. Formal kann für jede dieser Gitter das Fourierspektrum bestimmt werden. Die Autoren fanden nun, dass der Schwellenkontrast c_0 für ein Luminanzprofil dem Schwellenkontrast für die erste Harmonische des Spektrums eines Gitters entspricht. Für ein Rechteckgitter findet man die Reihe

$$u(x) = \frac{4}{\pi} \left(\sin \omega x + \frac{1}{3} \sin 3\omega x + \frac{1}{5} \sin 5\omega x + \dots \right), \quad (3.58)$$

(vergl. (A.49) auf Seite 163 im Anhang) sowie die Abbildung (A.3), Seite 164). Es folgt, dass die erste Harmonische eines Rechteckgitters mit dem Kontrast c gleich $4c/\pi$ ist. Die dritte Harmonische hat einen Kontrast von $4c/3\pi$, die fünfte hat einen Kontrast von $4c/5\pi$, etc. Alle Geradzahlgigen Harmonischen haben einen Kontrast von Null, – sie sind nicht im Spektrum enthalten. Soll das Rechteckgitter entdeckt werden, so wird es in erster Linie auf die erste Harmonische ankommen, und die Empfindlichkeit (sensitivity = $1/c_0$, c_0 der Schwellenkontrast) für das Rechteckgitter sollte um den Faktor $4/\pi$ größer sein als das für ein Sinusgitter mit der (Orts-)Frequenz der ersten Harmonischen des Rechteckgitters. Ein zweiter Test ergibt sich aus der Beobachtung, dass die nächsthöhere Harmonische die dreifache Ortsfrequenz der ersten hat, und mit nur $1/3$ der Amplitude der ersten eingeht. Die Stimulusdarbietung

Abbildung 3.21: Test der Campbell-Robson'schen Hypothese (1968): [A] (a) Kontrastsensitivität von Sinusgittern (kreisförmig) und Rechteckgittern (quadratisch) verschiedener Ortsfrequenz. (b) Der Quotient der Sensitivitäten wird mit $4/\pi = 1.273$ vorausgesagt (horizontale Linie). Die gestrichelte Linie ist die Vorhersage anhand eines einfachen Peak-Detectors. [B]



wurde nun so eingerichtet, dass alternierend ein Rechteckgitter und ein Sinusgitter, dessen Ortsfrequenz dem der ersten Harmonischen entsprach, gezeigt wurden, wobei der Kontrast des Sinusgitters um den Faktor 4π größer als der des Rechteckgitters war. Der Wechsel von einem Muster zum anderen ging nicht mit einem Wechsel der durchschnittlichen Helligkeit einher, und die beiden Gitter hatten die gleiche Phase. Die Versuchsperson konnte nun den Kontrast der beiden Gitter variieren, wobei aber das Verhältnis der Kontraste der beiden Gitter konstant blieb. Zunächst wurde der Kontrast als so gering vorgegeben, dass die beiden Gitter *nicht* unterschieden werden konnten, d.h. vom Rechteckgitter wurde nur die erste Harmonische wahrgenommen. Insbesondere *erhöhte* nun die Versuchsperson den Kontrast so lange, bis die beiden Gitter gerade als verschieden wahrgenommen wurden. Der ebenmerkliche Unterschied zwischen den beiden Gittern entsteht offenbar genau dann, wenn die dritte Harmonische gerade sichtbar wird.

Die Resultate sieht man in Abb.3.21, [B]. Die hellen Kreise rechts sind die Kontrastschwellen für die jeweilige Ortsfrequenz des Sinusgitters. Rückt man diese Punkte um drei Ortsfrequenzen nach links und dividiert die Kontrastschwellen durch 3, so erhält man die schwarzen Kreisflächen, – sie fallen auf die Kurve für die Kontrastschwellen des Rechteckgitters. Es lässt sich folgern, dass das Rechteck- und das Sinusgitter genau dann unterscheidbar werden, wenn die dritte Harmonische mit einer Amplitude erscheint, die der Sensitivität für diese Ortsfrequenz entspricht. Das Experiment wurde wiederholt, wobei das Rechteckgitter durch ein Sägezahn-gitter ersetzt wurde. Das Ergebnis ist das gleiche: die beiden Gitter werden unterscheidbar, wenn die jeweils nächsthöhere Ortsfrequenz mit hinreichender Amplitude

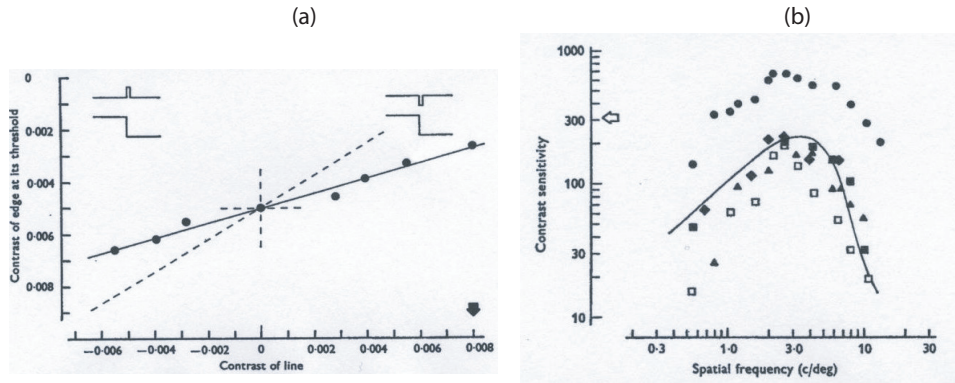
auftritt. Campbell und Robson diskutieren ihre Daten in Bezug auf ein Modell, dass einen einzigen Peak-Detektor für alle Ortsfrequenzen annimmt. Dieses Modell scheint mit den Daten nicht kompatibel zu sein. Als Alternative nehmen sie an, dass das visuelle System einzelne Ortsfrequenzfilter oder Ortsfrequenz"kanäle" enthält, die unabhängig voneinander angesprochen werden. Jeder dieser Kanäle ist durch eine eigene Systemfunktion charakterisiert. Die Kanäle können nicht vollständig unabhängig voneinander sein, da die Systemfunktionen keine Dirac-Delta-Funktionen sein können (das müssten sie sein, wären die Kanäle für jeweils nur eine Ortsfrequenz zuständig). Campbell und Robson argumentieren, dass die Bandbreite ungefähr ± 3 Ortsfrequenzen um die jeweils optimale Ortsfrequenz beträgt.

Blakemore und Campbell (1969) führten psychophysische Adaptationsexperimente durch: Sie bestimmten zunächst die Sensitivität für Sinusgitter mit verschiedenen Ortsfrequenzen und ließen die Versuchsperson dann an ein Gitter mit einer bestimmten Ortsfrequenz adaptieren. Die Adaptation bewirkt eine Erhöhung der Kontrastschwelle nicht nur für dieses Gitter, sondern auch für Gitter mit benachbarter Ortsfrequenz. Dieser Nachbarschaftsbereich beträgt ca. eine Oktave um die adaptierte Ortsfrequenz (Reduktion um eine halbe Amplitude); dies gilt für Ortsfrequenzen zwischen 3 c/deg und 14 c/deg. Für höhere Ortsfrequenzen wird die Bandbreite etwas geringer, für niedrigere als 3 c/deg bleibt das Maximum des Effekts bei 3 c/deg. Die Autoren äußern die Hypothese, dass das visuelle System selektiv für einzelnen Ortsfrequenzen ist. Campbell, Cooper und Enroth-Cugell (1969) lieferten Daten aus Einzelzellableitungen bei der Katze, die mit dieser Hypothese kompatibel sind. Maffei und Fiorentini (1973) lieferten weitere neurophysiologische Daten, die nahelegen, dass es insbesondere die *simple cells* sind, die auf spezifische Ortsfrequenzen reagieren und die also als Ortsfrequenzfilter interpretiert werden können. Sie fassen ihre Ergebnisse sowie die anderer Autoren zu der Hypothese zusammen, dass der visuelle Kortex bei der Mustererkennung eine Art Fourier-Analyse durchführt, d.h. dass die Musterrepräsentation schließlich auf der Basis einer Fourier-Synthese – als Synthese der Antworten der *simple cells* – geschieht.

3.4.5 Neuronale Kanäle: Kanten-, Linien- und Gitterdetektoren

Hubel und Wiesel (1962, 1968) hatten bereits gefunden, dass im visuellen Kortex der Katze und des Affen Neurone existieren, die nicht auf Linien oder Balken reagieren, aber auf Kanten, d.h. auf Grenzen zwischen hell und dunkel. Blakemore und Nachmias (1971) hatten gezeigt, dass die Adaption an bestimmte Ortsfrequenzen die Schwelle für Kanten in dem Masse erhöht, in dem die betreffende Ortsfrequenz im Spektrum dieser Reizmuster enthalten war. Tolhurst (1972) berichtete Experimente, bei denen die Versuchsperson an einen, wie Tolhurst ihn nennt, *left-right odd symmetry*-Stimulus adaptierte. Formal sind solche Stimuli durch $L(x) = -L(-x)$ definiert. Zeigt man nun einen Stimulus, dessen Asymmetrie die entgegengesetzte Polarität hat, so ist der Schwellenkontrast für ein solches Muster deutlich weniger erhöht als für ein Stimulmuster mit der gleichen Polarität. Tolhurst argumentiert, dieser Befund sei Evidenz für die Existenz von Kanälen, die spezifisch auf Kantenmuster bzw. auf Stimuli mit entsprechender Asymmetrie reagieren, es handele sich

Abbildung 3.22: Shapley und Tolhurst (1973): Test der Hypothese zur die Existenz von Kantendetektoren. (a) Unterschwellige Linie, (b) unterschwelliges Sinusgitter, jeweils superponiert auf einer Kante.



vermutlich um Kantendetektoren. Diese Hypothese wurde von

Shapley und Tolhurst (1973) im Rahmen von Superpositionsexperimenten weiter untersucht. Dazu wurden z. B. Linien oder Sinusgitter mit unterschwelligen Kontrasten einer Luminanzkante superponiert und die Kontrastschwelle für die Kante in Abhängigkeit vom Ort bzw. der Ortsfrequenz des Gitter bestimmt. Sie nahmen an, dass die angesprochenen Kanäle linear sind, so dass die Beziehung

$$c_t S_t + c_s S_s = k \quad (3.59)$$

gelten muß: hierin ist k eine Konstante, S_t ist die Sensitivität (= der Reziprokwert des Schwellenkontrasts) des Kanals für das Testmuster – hier: die Kante –, wenn es allein präsentiert wird, und S_s ist die Sensitivität des gleichen Kanals für das superponierte, unterschwellige Muster. c_t und c_s sind die Kontraste der beiden Muster in der Stimuluspräsentation. Gilt (3.59), so muß die Beziehung zwischen c_t und c_s linear sein. Die Abb. 3.22 (a) zeigt den experimentellen Befund: sicherlich ist die Beziehung linear, so dass die Grundannahme (3.59) für Schwellenkontraste akzeptiert werden kann. Der Kontrast c_s für die Linie wurde vorgegeben und der Kontrast c_t für die Kante wurde experimentell bestimmt. Der Pfeil auf der rechten Seite der x -Achse zeigt die Kontrastschwelle für das Linienelement an. (3.59) impliziert

$$c_t = \frac{k}{S_t} - c_s \frac{S_s}{S_t}, \quad (3.60)$$

d.h. die Steigung der *Kontrastinterrelationsfunktion* sollte durch den Quotienten S_s/S_t gegeben sein, vergl. die gestrichelte Gerade in Abb. 3.22 (a). Die Steigung der empirisch gefundenen Geraden weicht aber von dieser Vorhersage ab. Shapley und Tolhurst argumentieren, dass dieser Befund bedeutet, dass die Empfindlichkeit des visuellen Systems insgesamt für die dünne Linie größer ist als für die Kante.

Abbildung 3.23: Kulikowski & King-Smith (1973): Modell und erste Daten

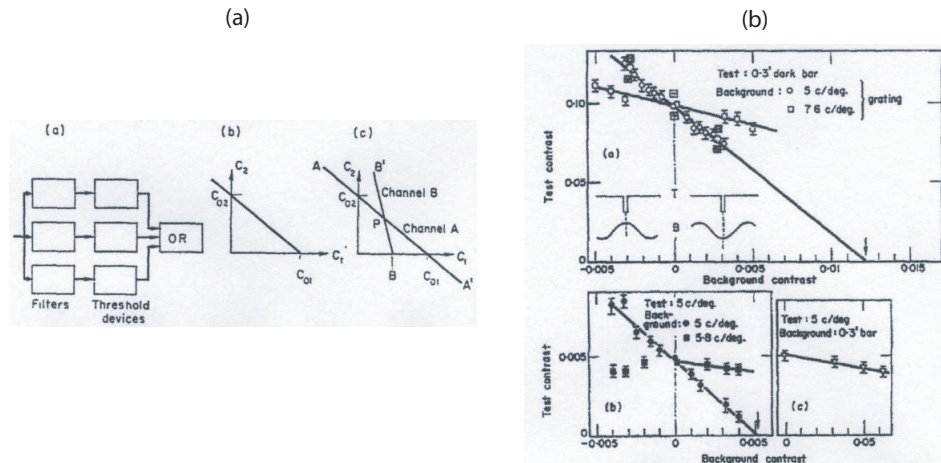


Abb. 3.22 (b) zeigt die Daten für superponierte Sinusgitter verschiedener Ortsfrequenz, wobei verschiedene "Kanten" muster als Testmuster dienten. Die "Diamonds" repräsentieren das Kantenmuster, dessen Luminanzprofil in (a) gezeigt wird. Die schwarzen Dreiecke zeigen die Resultate für Rechtecksgitter (also eine Folge von Kanten) für .4 c/deg, während die schwarzen und weißen Quadrate durch .53 c/deg charakterisiert waren (zwei verschiedene Messungen mit dem gleichen Testreiz). Die schwarzen Kreise sind die Sensitivitäten für einfache Sinusgitter (einfach, weil kein zusätzliches Testmuster gezeigt wurde). Die durchgezogene Linie zeigt das Amplitudenspektrum für eine Kante. Sie entspricht der Sensitivität für eine Kante in Abhängigkeit von der Ortsfrequenz. Auf die Diskussion der Hypothese soll hier nicht in allen Details eingegangen werden, aber die Zusammenfassung, die Shapley zu dieser Diskussion geben, soll direkt zitiert werden:

... , a picture emerges of a visual system thick with channels: some narrow-band width grating channels, and some broad-band, localized feature detectors. For detection tasks only one or a few channels are presumably operating; but in everyday suprathreshold vision all will be activated. What other kinds of feature detector exist and how the various types of detector interact are important questions that remain to be answered.

Kulikowski & King-Smith (1973) gehen ebenfalls die Frage nach den möglichen Kanälen an, sind dabei aber noch systematischer als Shapley und Tolhurst. Das von ihnen zugrundegelegte Modell ist schematisch in Abb. 3.23 (a) abgebildet. Demnach existiert eine Menge von voneinander unabhängigen Filtern, die über eine ODER-Verbindung in die Entscheidung der Versuchsperson eingehen: diese gibt eine "ja"-Antwort (d.h. es ist ein Stimulus gezeigt worden), wenn der maximal aktivierte Kanal auch hinreichend aktiviert wurde. Das Modell ist insofern deterministisch, als keine W-Summation zwischen Kanälen angenommen wird – man kann sagen, dass die tatsächliche Annahme ist, dass das Rauschen in den Antworten der verschiede-

nen Filter so gering ist, dass W-Summationseffekte nicht auftreten; eine Alternative Annahme ist, dass das Rauschen im System der Kanäle insgesamt zu vernachlässigen ist und nur am Ort der Entscheidung eine Rolle spielt. (b) zeigt die Kontrast-Interrelationsfunktion, die formal wie in (3.59) definiert ist: gezeigt wird ein Stimulus, für den ein Kanal maximal empfindlich ist, und superponiert wird ein zweiter Stimulus – etwa ein Sinusgitter – gezeigt. Für vorgegebenen Kontrast des einen (Teil-)Musters ergibt sich der Kontrast des anderen Muster, damit der Stimulus gerade entdeckt wird. Die Beziehung zwischen den beiden Kontrasten sollte, wie in (3.60), linear sein. Es wird also angenommen, dass der für den "eigentlichen" Stimulus zuständige Kanal auch für den perturbierenden, superponierten Stimulus empfindlich ist. (c) zeigt, was geschieht, wenn der Kontrast für das perturbierende Muster hinreichend groß und der für den "eigentlichen" Stimulus dementsprechend klein wird, dass nun der Kanal für das perturbierende Muster zum entdeckenden Kanal wird, der durch das andere Muster nur mitaktiviert wird: Die Kontrast-Interrelation (I) geht in die Kontrastinterrelation (II) über, die durch eine Gerade mit anderen Parametern charakterisiert ist. Formal läßt sich diese Beziehung wie folgt zeigen: Es sei c_0 der Schwellenkontrast für einen Stimulus (Kante, Linie, Gitter). Für einen etwas von c_0 abweichenden Kontrast c weicht die Antwort des Kanals von Schwellenaktivität so ab, dass die Reaktion durch eine lineare Funktion beschrieben werden kann (Linearität ist letztlich nur eine Approximation, also kann man keinen beliebigen Kontrast annehmen). Kulikowski et al. betrachten nun die normierte Antwort

$$R = \frac{c}{c_0}. \quad (3.61)$$

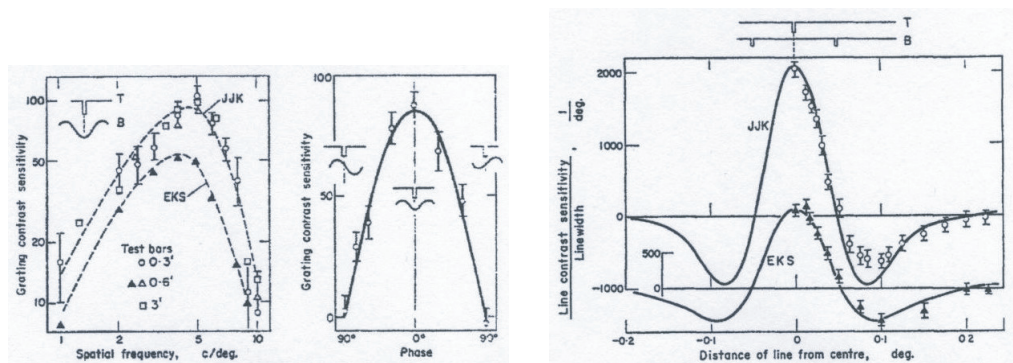
Die normierte Antwort ist also proportional zur Sensitivität $1/c_0$ des Kanals, und für $c = c_0$ ist $R = 1$. Für eine Superposition zweier Muster soll nun gelten, dass die Gesamtaktivität gleich der Summe der Teilaktivitäten ist. Dies folgt aus der Linearitätsannahme. Es sei nun c_{01} der Schwellenkontrast des Kanals für das erste Muster, und c_{02} der Schwellenkontrast des gleichen Kanals für das zweite Muster. Der durch die Superposition definierte Stimulus wird gerade wahrgenommen, wenn

$$R = \frac{c_1}{c_{01}} + \frac{c_2}{c_{02}} = 1. \quad (3.62)$$

Die Beziehung zwischen c_1 und c_2 ist demnach eine Gerade in einem durch C_1 und C_2 definierten Koordinatensystem. Für $c_2 = 0$ folgt $c_1/c_{01} = 1$, d.h. $c_1 = c_{01}$, und analog folgt für $c_1 = 0$ die Beziehung $c_2 = c_{02}$; die Gerade schneidet also die x - bzw. y -Achse in c_{01} bzw. c_{02} . Man bemerke, dass c_{01} und c_{02} die Schwellenkontraste der beiden Stimuli *einen und denselben Kanal* sind. Besteht also der Gesamtstimulus aus der Superposition einer Kante und eines Sinusgitter, so läßt sich im Prinzip die Sensitivität des Kantendetektors für ein Sinusgitter mit der gegebenen Ortsfrequenz bestimmen, und die Sensitivität des Gitterdetektors – für die gegene Ortsfrequenz – für eine Kante.

Abb. 3.23 (b) zeigt die ersten Resultate. Eine dünne Linie wurde einem Sinusgitter mit entweder der Ortsfrequenz 5 c/deg oder 7.6 c/deg superponiert. Die Gitter konnten entweder in Sinus- oder Kosinusphase relativ zur Linie sein; die Sinusphase impliziert dann einen "negativen" Kontrast. Die Kontrastinterrelationen für die

Abbildung 3.24: Kulikowski & King-Smith (1973): Gittersensitivität von Detektoren für Linien von .3', .6' und 3' Breite (Minuten/Schwinkel), für zwei Versuchspersonen JJK und EKS. Außerdem wird die Empfindlichkeit des Liniendetektors Gitter in verschiedenen Phasenlagen gezeigt. Die Sensitivitätsfunktion entspricht den Funktionen, die andere Autoren für die Sensitivität von Ganglienzellen gefunden haben (Enroth-Cugell und Robson (1966)), und für kortikale Zellen (Campbell, Cooper und Enroth-Cugell (1969)). (b) Die Empfindlichkeit des Liniendetektors für flankierende Linien (*Line-spread-function*).



beiden Gitter + Linie Stimuli sind sicherlich linear, d.h. die entdeckenden Kanäle sind sicherlich im Schwellenbereich durch lineare Systeme zu approximieren. Die Pfeile auf der x -Achse zeigen auf die Schwellenkontraste des Liniendetektors für Sinusgitter; dies sind also hypothetische Kontraste, für die der Liniendetektor durch ein Sinusgitter bis zum Schwellenwert aktiviert wird.

Abb. 3.24 (a) zeigt die Empfindlichkeit von Liniendetektoren für Sinusgitter, und die Empfindlichkeit für Sinusgitter in Abhängigkeit von der Phase. Für die Kosinusphase ist die Empfindlichkeit maximal, für die Sinusphase ist sie minimal. In (b) wird die Empfindlichkeit des Liniendetektors für flankierende Linien in Abhängigkeit von der Distanz der flankierenden Linien gezeigt. Die Kurve zeigt eine von Null verschiedene positive Empfindlichkeit für benachbarte Linien (bis ca. .005' Bogenminuten), die dann in eine inhibitorische Reaktion übergeht, die für größer werdende Distanz gegen Null geht. Die durchgezogene Linie entspricht der Differenz zweier Gauß-Funktionen; die Funktion, die die inhibitorischen Effekte repräsentiert, hat den größeren Varianzparameter.

Abadi & Kulikowski (1973) haben die Bandbreite der Ortsfrequenzkanäle im Paradigma von Superpositionsexperimenten weiter untersucht. Teststimulus war ein Sinusgitter von 5 c/deg, dass entweder einem Sinusgitter gleicher Frequenz, oder einem Rechteckgitter superponiert wurde. Das Rechteckgitter hatte eine Ortsfrequenz von entweder ebenfalls 5 c/deg, oder 1.67 c/deg, oder 1 c/deg. Die Reihe für ein

Rechtecksgitter hat die Form

$$s(x) = \frac{4}{\pi} \left(\sin(\omega x) + \frac{1}{3} \sin(3\omega x) + \frac{1}{5} \sin(5\omega x) + \dots \right) \quad (3.63)$$

Die Antwort auf die Superposition hat die Form

$$g = c_T h_T + c_B h_B,$$

wobei c_T der Kontrast des Teststimulus und T_B der Hintergrundmusters ist. h_T und h_B sind die Einheitsantworten (Einheit bezieht sich hier auf einen Kontrast gleich 1) des entdeckenden Kanals auf den Test- bzw. den Hintergrundreiz. Für $g = k$ wird das Muster mit der Wahrscheinlichkeit p_0 – etwa $p_0 = .5$ – entdeckt (zeitliche Peak-Detektion). Absorbiert man den Faktor $1/k$ in h_T und h_B , so ergibt sich die Kontrastinterrelation

$$c_T = \frac{1}{h_T} - c_B \frac{h_B}{h_T}. \quad (3.64)$$

Die Steigung dieser Geraden ist durch das Verhältnis h_B/h_T der Einheitsantworten gegeben. Abadi et al.s Vermutung war nun, dass das Hintergrundmuster nur durch diejenige Fourier-Komponente wirkt, die der Orstfrequenz des Stimulsmusters entspricht. Für das 5 c/deg-Muster wäre dies die erste Komponente, die mit einem Faktor $4/\pi$ in das Hintergrundmuster eingeht. Es ist

$$\frac{4/\pi}{1} = \frac{4}{\pi} = 1.27324 \approx 1.27, \quad \frac{4}{3\pi} \approx .42, \quad \frac{4}{5\pi} \approx .255$$

Abadi et al. fanden Steigungen von $1.289 \pm .089$, $.402 \pm .027$ und $.195 \pm .036$, was den Vorhersagen im Rahmen experimenteller Fehler entspricht.

3.4.6 Neuronale Kanäle als Matched Filter

Wenn von Kanten-, Linien und Gitterkanälen die Rede ist, so wird damit ausgesagt, dass diese Kanäle maximal empfindlich sind für eben diese Kanten, Linien und Gitter. Linien können eine bestimmte Breite haben, Kanten müssen nicht notwendig durch eine Schrittfunktion definiert sein, sondern können wie eine Rampe oder sigmoidal verlaufen. Die Frage ist, ob es Kanäle für jede dieser Varianten gibt oder nicht, und wenn ja, für welche Musterkomponenten es überhaupt Kanäle gibt. Die Superpositionsexperimente, die von Shapley und Tolurst und insbesondere von Kulikowski und King-Smith durchgeführt wurden legen nahe, dass die Hypothese, dass ein spezieller Kanal für irgendein vorgegebenes Muster (zunächst noch 1-dimensional definiert) ebenfalls über Superpositionsexperimente diskutiert werden kann. Die Systemfunktion eines Kanals, der für ein bestimmtes Muster maximal reagiert, lässt sich leicht ableiten. Es sei H die Systemfunktion eines Kanals und S sei die Fouriertransformierte eines Signals (Stimulus). Die Antwort an der retinalen Position x dieses Kanals ist dann durch

$$g(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} H(\omega) S(\omega) e^{i\omega x} d\omega$$

gegeben. In Abschnitt 1.2.11, Gleichung (1.181) wurde die Beziehung zwischen S und H aufgeführt, die an einem Punkt x_0 die Antwort maximiert. Für $x_0 = 0$ wird der Stimulus foveal präsentiert und der Filter für foveal präsentierte Stimuli der entsprechenden Art reagiert. Die Relation (1.181) erlaubt es, die Systemfunktion des angesprochenen Kanals vorherzusagen, – vorausgesetzt, die zeitliche Peak-Detection-Annahme gilt und die Versuchsperson orientiert sich tatsächlich an der Aktivität in x_0 , ob nun $x_0 = 0$ oder $x_0 \neq 0$ gilt.

Nun gehen die Übertragungsfunktion der Linse, der Retina und des *corpus geniculatum laterale* in die Filtercharakteristika mit ein. Es ist denkbar, dass diese Stationen der Musterverarbeitung wie ein Vorfilter wirken und derusterspezifische Kanal auf die Antwort dieses Vorfilters auf den Stimulus arbeitet. Wird zusammenfassend dieser Vorfilter durch eine Systemfunktion $H_0(\omega)$ repräsentiert und hat der Stimulus die Fouriertransformierte $S(\omega)$, so ist die Fouriertransformierte der Antwort des Vorfilters durch

$$G_0(\omega) = H_0(\omega)S(\omega) \quad (3.65)$$

gegeben. Der entdeckende Filter muß dann an diese Antwort des Vorfilter, der das Signal für den entdeckenden Filter ist, angepaßt sein, d.h. seine Systemfunktion muß die zu $G_0(\omega)$ konjugiert komplexe Funktion sein. Die Antwort des angepaßten Filters für irgendein x ist durch

$$g_s(x, c) = \frac{c}{2\pi} \int_{-\infty}^{\infty} H_0(\omega)(S(\omega)[H_0(\omega)S(\omega)]^* e^{i\omega x} d\omega \quad (3.66)$$

gegeben. Es werden nun die folgenden Annahmen gemacht:

- A1: Für einen foveal präsentierten Stimulus wird die Entdeckensantwort durch den Wert von g_s in $x = x_0 = 0$ bestimmt,
- A2: Der Stimulus wird durch den angepaßten Filter entdeckt, wenn $g_s(0) = k$, k eine Konstante.
- A3: Der Vorfilter H_0 ist eine gerade Funktion von ω , d.h. $H_0(\omega) = H_0(-\omega)$.

Für die Antwort in $x = 0$ erhält man dann aus (3.66)

$$g_s(0, c) = \frac{c}{2\pi} \int_{-\infty}^{\infty} H_0(\omega)S(\omega)[H_0(\omega)S(\omega)]^* d\omega = \frac{c}{2\pi} \int_{-\infty}^{\infty} |H_0(\omega)S(\omega)|^2 d\omega \quad (3.67)$$

gleichbedeutend ist. Da der Experimentator den Stimulus $s(x)$ vorgibt, ist $S(\omega)$ als Fouriertransformierte von s bekannt. Gelingt es jetzt auch noch $H_0(\omega)$ abzuschätzen, kann man im Prinzip den Wert von $g_s(x)$ für beliebige x berechnen. Das Problem ist nun aber, H_0 zu schätzen, denn man muß damit rechnen, dass für jeden Stimulus s , den man wählt, die Beziehung (3.67) besteht, – die Aufgabe ist aber, anzuschätzen, ob überhaupt ein Filter für s existiert, der an s angepaßt ist.

Um diese Aufgabe anzugehen, wird zunächst die Systemfunktion für einen Kanal, der maximal im Sinne des Matched Filters auf ein Signal s reagiert, angeschrieben.

Nach (3.65) ist die Fouriertransformierte der Antwort des Vorfilters auf s durch $G_0 = H_0 S$ gegeben, und der nun folgende Matched Filter hat die Systemfunktion $(H_0 S)^*$. Der Vorfilter und der Matched Filter sind zwei in Reihe geschaltete Filter: zunächst H_0 , dann $(H_0 S)^*$, und diese Folge von Filtern hat dann die Systemfunktion

$$H_s(\omega) = H_0(\omega)(H_0(\omega)S(\omega))^* = |H_0(\omega)|^2 S^*(\omega). \quad (3.68)$$

Man kann nun die Fouriertransformierte F_b der Antwort auf einen beliebigen Stimulus s_b auf den durch H_s charakterisierten Kanal bestimmen:

$$F_b(\omega) = \begin{cases} |H_0(\omega)|^2 S^*(\omega) S_b(\omega), & s_b \neq s \\ |H_0(\omega)|^2 |S(\omega)|^2, & s_b = s \end{cases} \quad (3.69)$$

Die Hypothese des Entdeckens durch einen Matched Filter für einen vorgegebenen Stimulus s kann wieder über ein Superpositionsexperiment getestet werden. Dazu wird der Stimulus auf einem Sinus- oder Kosinusgitter mit Ortsfrequenz $f = \omega/2\pi$ präsentiert und der Schwellenkontrast $c_0(f)$ für s in Abhängigkeit von f bestimmt, wobei f einen hinreichend großen Bereich durchläuft. Es läßt sich nun zeigen, dass die resultierende Schwellenkurve $c_0(f)$ versus f proportional zum Realteil der Systemfunktion des entdeckenden Kanals ist, wenn das Hintergrundgitter in Kosinusphase ist, und proportional zum Imaginarteil der Systemfunktion, wenn das Hintergrundgitter in Sinusphase ist (jeweils relativ zu $x = 0$).

Um diese Behauptung einzusehen, wird zunächst die den Gesamtstimulus definierende Superposition angeschrieben:

$$\sigma(x) = c_s s(x) + \begin{cases} c_B \sin(\omega_0 x), & \text{Sinusphase} \\ c_B \cos(\omega_0 x), & \text{Kosinusphase} \end{cases} \quad (3.70)$$

Die Fouriertransformierte von $\sigma(x)$ ist dann

$$\Sigma(\omega) = c_s S(\omega) + \begin{cases} -c_B i\pi [\delta(\omega - \omega_0) - \delta(\omega + \omega_0)] \\ c_B \pi [\delta(\omega - \omega_0) + \delta(\omega + \omega_0)] \end{cases} \quad (3.71)$$

Die Gleichung (3.69) liefert nun die Fouriertransformierte der Antwort des hypothetischen Angepaßten Filters auf das Stimulismuster $\sigma(x)$. Denn unter Berücksichtigung von (3.68) erhält man zunächst

$$F_\sigma(\omega) = H_s(\omega) \Sigma(\omega),$$

und dann

$$F_\sigma(\omega) = c_s |H_0(\omega)|^2 |S(\omega)|^2 + \begin{cases} -c_B i\pi |H_0(\omega)|^2 [\delta(\omega - \omega_0) - \delta(\omega + \omega_0)] \\ c_B \pi |H_0(\omega)|^2 [\delta(\omega - \omega_0) + \delta(\omega + \omega_0)] \end{cases} \quad (3.72)$$

Die Antwort des durch H_s definierten Filters in $x = x_0 = 0$ ist dann

$$g_s(0) = \frac{1}{2\pi} \int_{-\infty}^{\infty} H_s(\omega) \Sigma(\omega) d\omega.$$

Setzt man für Σ den in (3.72) gegebenen Ausdruck ein, so erhält man

$$g_s(0) = \frac{c_s}{2\pi} \int_{-\infty}^{\infty} |H_0(\omega)|^2 |S(\omega)|^2 d\omega + c_B R_B(\omega) \quad (3.73)$$

mit

$$R_B(\omega) = \begin{cases} -\frac{i}{2}(|H_0(\omega_0)|^2 S^*(\omega_0) - |H_0(-\omega_0)|^2 S^*(-\omega_0)) \\ \frac{1}{2}(|H_0(\omega_0)|^2 S^*(\omega_0) + |H_0(-\omega_0)|^2 S^*(-\omega_0)) \end{cases} \quad (3.74)$$

Weiter sei

$$S(\omega) = S_1(\omega) + iS_2(\omega), \quad i = \sqrt{-1} \quad (3.75)$$

d.h. S_1 sei der Real-, und S_2 sei der Imaginärteil von S . Dann folgt wegen der Annahme A3:

$$|H_0(\omega_0)|^2 S^*(\omega_0) - |H_0(-\omega_0)|^2 S^*(-\omega_0) = |H_0(\omega_0)|^2 (S_1(\omega_0) + iS_2(\omega_0) - S_1(-\omega_0) - iS_2(-\omega_0)).$$

Nun ist der Realteil einer Fouriertransformierten stets eine gerade Funktion von ω bzw. ω_0 , und der Imaginärteil ist eine ungerade Funktion, so dass

$$S_1(\omega_0) - S_1(-\omega_0) = 0, \quad iS_2(\omega_0) - iS_2(-\omega_0) = 2iS_2(\omega_0).$$

Damit hat man für R_B im Falle der Sinusphase des Hintergrundgitters $R_B(\omega_0) = \pi S_2(\omega_0)$. Auf analoge Weise findet man für die Kosinusphase des Hintergrundgitters $R_B(\omega_0) = \pi S_1(\omega_0)$; zusammenfassend erhält man das Resultat

$$R_B(\omega_0) = \begin{cases} \pi |H_0(\omega_0)|^2 S_2(\omega_0), & \text{Sinusphase} \\ \pi |H_0(\omega_0)|^2 S_1(\omega_0), & \text{Kosinusphase} \end{cases} \quad (3.76)$$

Es werde nun zunächst der Fall betrachtet, dass das Testmuster ohne ein unterlegtes Gitter präsentiert wird und der Schwellenkontrast für dieses Muster gleich c_{0s} ist. nach A2 wird das Muster entdeckt, wenn g_{0s} – die Antwort des entdeckenden Kanals, wenn kein Hintergrundgitter präsentiert wird – gleich k ist, k eine bestimmte Konstante, wenn also

$$\frac{c_{0s}}{2\pi} \int_{-\infty}^{\infty} |H_0(\omega) S(\omega)|^2 d\omega = k \quad (3.77)$$

gilt. Dann folgt

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} |H_0(\omega) S(\omega)|^2 d\omega = \frac{k}{c_{0s}}, \quad (3.78)$$

Damit folgt aus (3.73) für das Entdecken

$$g_s(0) = k \frac{c_s}{c_{0s}} + c_B R_B(\omega_0) = k. \quad (3.79)$$

Für gegebenen Wert von ω_0 , der Ortsfrequenz des Hintergrundmusters, ist die Beziehung zwischen c_s und c_B offenbar linear. Insbesondere erhält man für R_B den Ausdruck

$$R_B(\omega) = \frac{1}{c_B} \left(k - k \frac{c_s}{c_{0s}} \right), \quad (3.80)$$

bzw.

$$\frac{1}{k} \hat{R}_B(\omega_0) = \frac{c_{0s} - c_s}{c_{0s} c_s}. \quad (3.81)$$

Die rechte Seite enthält nur empirisch bestimmte Größen, während links ein von den bekannten Größen $S_1(\omega_0)$ oder $S_2(\omega_0)$ und der unbekannten Größe $(\pi/k)|H_0(\omega_0)|^2$ abhängender Term steht. Es sei nun nach (3.76),

$$R_{B,\sin}(\omega_0) = \pi |H_0(\omega_0)|^2 S_2(\omega_0), \quad R_{B,\cos} = \pi |H_0(\omega_0)|^2 S_1(\omega_0). \quad (3.82)$$

Für die Größen $R_{B,\sin}(\omega_0)$, und $R_{B,\cos}$ liegen die empirischen Abschätzungen (3.81) vor. Dividiert man sie durch $S_2(\omega_0)$ bzw. durch $S_1(\omega_0)$, so erhält man die Schätzungen

$$\frac{(1/k) \hat{R}_{B,\sin}(\omega)}{S_2(\omega)} \approx |H_0(\omega_0)|^2, \quad \frac{(1/k) \hat{R}_{B,\cos}(\omega)}{S_1(\omega)} \approx |H_0(\omega_0)|^2. \quad (3.83)$$

Man erhält also zwei Schätzungen für $|H_0(\omega_0)|^2$. Man kann auch die vorhandene Information über $|H_0(\omega_0)|^2$ gleich zusammenfassen, um eine möglicherweise bessere Schätzung zu bekommen. Es ist ja

$$\begin{aligned} \sqrt{R_{B,\sin}(\omega_0)^2 + R_{B,\cos}(\omega_0)^2} &= \pi |H_0(\omega_0)|^2 \sqrt{S_1^2(\omega_0) + S_2^2(\omega_0)}, \\ \sqrt{R_{B,\sin}(\omega_0)^2 + R_{B,\cos}(\omega_0)^2} &= \pi |H_0(\omega_0)|^2 |S(\omega_0)|^2. \end{aligned} \quad (3.84)$$

$|S(\omega_0)|^2$ ist aber bekannt. Schreibt man nun $c_{s,\sin}$, wenn c_s für das Hintergrundgitter in Sinusphase bestimmt wurde, und $c_{s,\cos}$, wenn c_s für das Hintergrundgitter in Kosinusphase bestimmt wurde, so erhält man die Abschätzung

$$\frac{1}{\pi} \sqrt{\left(\frac{c_{0s} - c_{s,\sin}}{c_{0s} c_{s,\sin}} \right)^2 + \left(\frac{c_{0s} - c_{s,\cos}}{c_{0s} c_{s,\cos}} \right)^2} \approx |H_0(\omega_0)|^2. \quad (3.85)$$

Man muß also einerseits den Schwellenkontrast c_{0s} für das Stimulsmuster s ohne Hintergrund bestimmen, und dann die Schwellenkontraste $c_{s,\sin}$ und $c_{s,\cos}$ für das Stimulsmuster mit Hintergrund, für eine Reihe von ω_0 -Werten. Dann lassen sich die $|H_0(\omega_0)|^2$ -Werte über (3.85) abschätzen. Aufgrund von (3.81) und (3.76) lassen sich jetzt die $R_b(\omega_0)$ -Werte für die verschiedenen ω_0 und die beiden möglichen Phasen des Hintergrundmusters voraussagen. Damit hat man einen Test der Hypothese, dass für ein vorgegebenes Muster s ein angepasster Filter existiert.

Hauske, Wolff und Lupp (1976) haben als Erste einen Versuch der beschriebenen Art durchgeführt. Abb. 3.25 zeigt die von den Autoren verwendeten Stimuli. Während Tolhurst (1972) und Shapley und Tolhurst (1973) nur die Hypothese der Existenz von Kantendetektoren geprüft haben und auch Kulikowski und King-Smith (1973) kaum über diese Klassen von Stimuli hinausgingen, wird bei Hauske et al.

Abbildung 3.25: Luminanzprofile der Stimuli in der Untersuchung von Hauske, Wolff und Lupp (1976). (a) eine tiefpassgefilterte Kante sowie eine Schrittfunktion, (b) eine als Rechteckprofil definierte senkrechte Linie, sowie eine doppeltdifferenzierte Linie, (c) helle Linien mit zwei unmittelbar benachbarten dunklen Streifen, (c) zwei unmittelbar benachbarte Halbzyklen eines Kosinusgitters ($|\cos(\omega_0 x)|$), sowie eine Exponentialfunktion.

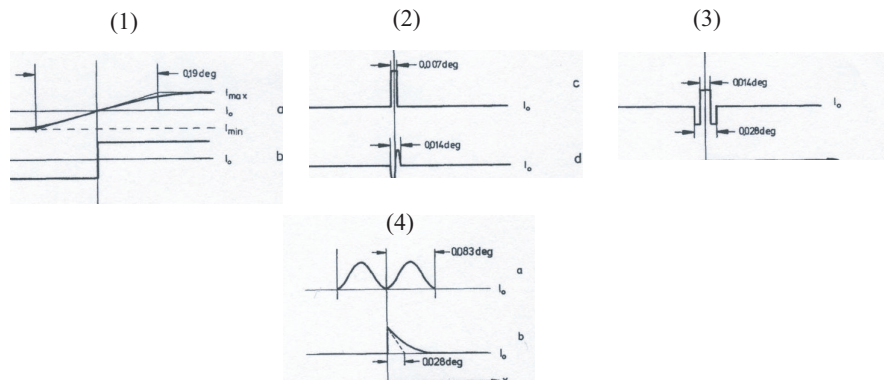
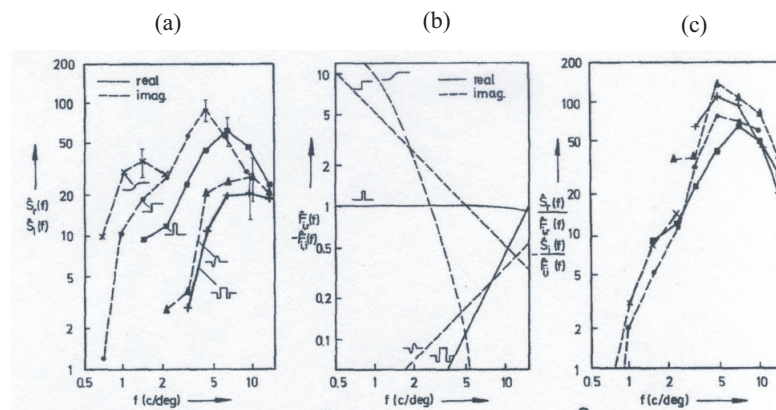


Abbildung 3.26: Resultate für die Stimuli in Abb. 3.25 (1) - (3). (a) Sensitivitätsfunktionen für die indizierten Stimuli, (b) Spektren (Real- und Imaginärteil der Muster, (c) Quotient des geschätzten Teils des Spektrum und des entsprechenden Teils des Stimulusspektrums; der Quotient ist eine Schätzung für $|H_0(\omega_0)|^2$, gemäß (3.83)).



eine größere Vielfalt von Stimuli betrachtet. Der eigentliche Test der Matched-Filter-Hypothese findet sich in Abb. 3.26 (c). Die Kurven sind Schätzungen von $|H_0(\omega_0)|^2$ und sollten also unabhängig vom verwendeten Stimulussmuster im Prinzip identisch sein. Stellt man die unvermeidlichen Meßfehler in Rechnung, so kann man sagen, dass die Hypothese mit den Daten kompatibel ist.

Abb. 3.27 zeigt die Daten für die Stimuli in Abb. 3.25 (4). Die Interpretation der

Abbildung 3.27: Resultate für die Stimuli in Abb. 3.25 (4). (a) Sensitivitätsfunktionen für die indizierten Stimuli, (b) Spektren (Real- und Imaginärteil der Muster, (c) Quotient des geschätzten Teils des Spektrum und des entsprechenden Teils des Stimulusspektrums; der Quotient ist wieder eine Schätzung für $|H_0(\omega_0)|^2$, gemäß (3.83)).

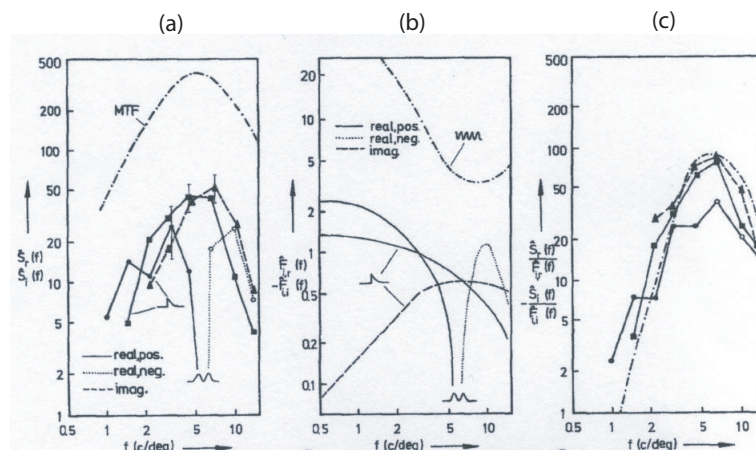


Abbildung korrespondiert zu der für die Abb. 3.26. In (a) wird noch die Schätzung der Modulationsübertragungsfunktion (MTF) für einfache sinuoidale Gitter gezeigt. Die Schätzungen für $|H_0(\omega_0)|^2$ für den Doppel-Kosinus-Stimulus liegen etwas niedriger als die für die aus den Real- und Imaginärteilen das anderen Musters (Exponentialfunktion). Man kann sagen, dass die Daten der Hypothese insgesamt nicht widersprechen. Die Hypothese wird allerdings kritisiert: es sei unplausibel, dass das visuelle System über eine große Anzahl von musterspezifischen Filtern verfüge. Denn letztlich sind hier nur einige Spezialfälle von Mustern untersucht worden, für die zudem gilt, dass der Punkt maximaler Antwort an der Stelle $x_0 = 0$ liegt. Stimuluster sind aber eigentlich 2-dimensional, – der Fokus auf 1-dimensionale Muster ergibt sich ja nur aus dem Wunsch, zunächst einmal einfache Muster zu untersuchen. Soll die Hypothese aber allgemeine Gültigkeit haben, so muß auch geprüft werden, ob für andere Muster der Ort maximaler Antwort auch an Punkten $(x_0, y_0) \neq (0, 0)$ liegen könnte. Eine weitere Frage ist die nach dem Entdecken komplexerer Muster, die aus Mustern der hier betrachteten Art zusammengesetzt sein können. Die Vermutung, dass das visuelle System nur über wenige *Feature Detectors* wie etwa Ortsfrequenzfilter verfügt, aus deren kombinierter Aktivierung sich die neuronale Repräsentation des komplexen Musters ergibt, ist sicherlich sehr viel plausibler als die Hypothese, dass das visuelle System über eine Unzahl aspektspezifischer Detektoren verfügt.

Folgt man dieser Kritik, so wird man auf die Gegenfrage geführt: kann man etwa nur mit Ortsfrequenzfiltern die Ergebnisse von Hauske et al. erklären? Insbesondere Graham (1989) argumentierte, dass die Daten der Art, wie Hauske et al. sie vorstellen, durch das Prinzip der Wahrscheinlichkeitssummation zwischen Ortskanälen erklärt werden können. Die Anpassung eines solchen W-Summationsmodells an die

Daten ist aber anscheinend nie durchgeführt worden. Die Argumentation, wie sie von Graham vorgebracht wurde, läßt aber die Möglichkeit außer Acht, dass die von Hauske et al. gefundenen Filter gar nicht für sich existieren, sondern durch Prozesse des perzeptuellen Lernens erst während des Experimentierens entstehen; bereits Hauske et al. merkten – eher beiläufig – an, dass ihre Befunde als eine Art Adaptationsprozess des visuellen Systems verstanden werden können. In der Tat läßt sich zeigen, dass die Wirkung der Hebb-Regel impliziert, dass Neurone, deren rezeptives Feld sich nach Maßgabe der Hebb-Regel modifiziert, bei mehrfacher Darbietung ein und desselben Stimulus gegen einen Filter konvergieren, der für den Musteraspekt, der das rezeptive Feld des Neurons überdeckt, zu einem Matched Filter wird (Nachtigall, 1991). Man würde dementsprechend vorhersagen, dass (i) Matched Filter nur für hinreichend kleine Stimuli beobachtet werden können, und die Hypothese des Entdeckens durch Matched Filter modifiziert werden muß: der Stimulus wird nicht durch einen Matched Filter entdeckt, der auf das Gesamtmuster reagiert, sondern durch denjenigen, der durch einen Musteraspekt maximal relativ zur Erregung der übrigen aktiviert wird. Dieses Modell wird weiter unten ausführlich vorgestellt. Hier sollen zunächst weitere Daten vorgestellt werden, die nahelegen, dass in der Tat Matched Filter existieren können, möglicherweise als Resultat eines Lernprozesses.

Meinhardt (1995) hat in einer Reihe von Versuchen der Frage nachgegangen, ob Effekte der Wahrscheinlichkeitssummation die Daten erklären können, die mit der Matched-Filter-Hypothese kompatibel sind.

hier kapitel über W-Summation zwischen Ortsfrequenzkanälen und spatiale
W-Summation

Anhang A

Anhang: Fourier- und Laplace-Transformationen

Die Diskussion des Verhaltens linearer dynamischer Systeme mit konstanten Koeffizienten wird oft vereinfacht oder veranschaulicht, wenn man von der Tatsache Gebrauch macht, dass viele Funktionen als additive Überlagerung von Sinus- und Kosinusfunktionen dargestellt werden können. Diese Darstellung geht auf Fourier zurück, und man spricht von der Fourier-Analyse bzw. Fourier-Transformation. Eine Verallgemeinerung wurde von Laplace vorgeschlagen. In diesem Abschnitt sollen einige Resultate der Theorie der Fourier- bzw. Laplace-Transformation angegeben werden. Auf Beweise oder Herleitungen muß dabei zum Teil verzichtet werden. Der interessierte Leser sei dazu auf die Literatur verwiesen (z.B. Bracewell (1978), Papoulis (1962)). Eine insbesondere für das Selbststudium sehr gut geeignete Einführung in die Theorie der Fourier- und Laplace-Transformation gibt Hsu (1967).

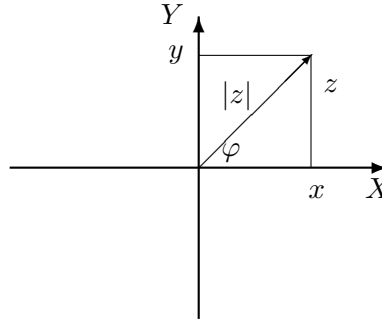
Voraussetzung für die Fourier- und Laplace-Analyse ist der Begriff der komplexen Zahl. Es werden zunächst einige der wichtigsten Eigenschaften komplexer Zahlen zusammengestellt.

A.1 Komplexe Zahlen

Gegeben sei die Gleichung $x^2 + 1 = 0$; gesucht ist eine Lösung für x . Da $x^2 = -1$ findet man $x_{1,2} = \pm\sqrt{-1}$. Diese Zahl ist offenbar nicht "reell". Gleichwohl erweist es sich als nützlich, mit ihr zu rechnen. Dazu führt man das Symbol $i = \sqrt{-1}$ ein. Aus der Definition der Wurzel folgt, dass dann $i^2 = -1$. Dementsprechend kann man weitere Potenzen i^k bilden, $k \in \mathbb{N}$, $\mathbb{N}$ die Menge der natürlichen Zahlen 1, 2, 3, So ist $i^3 = i \cdot i^2 = -i$, $i^4 = i^2 \cdot i^2 = -1 \cdot -1 = 1$, $i^5 = i \cdot i^4 = i \cdot 1 = i$, etc.

Man kann nun die Menge $\mathbb{R}$ der reellen Zahlen erweitern zur Menge $\mathbb{C}$ der komplexen Zahlen $z = x + iy \in \mathbb{C}$ mit $x \in \mathbb{R}$, $y \in \mathbb{R}$. Damit ist z durch das Zahlenpaar (x, y) bestimmt. Formal entspricht diesem Paar ein Vektor, und damit ist z durch graphisch durch einen Vektor im *Argand Diagramm* darstellbar: Die X -Achse bildet den Realteil x der komplexen Zahl ab, und die Y -Achse den Imaginärteil y (man

Abbildung A.1: Argand Diagramm für $z = x + iy$



bemerke, dass y , nicht iy abgetragen wird). $|z|$ ist die Länge des Vektors z ; $|z|$ wird auch als *Betrag* der komplexen Zahl bezeichnet.

Zwei komplexe Zahlen $z = x + iy$ und $z' = x' + iy'$ sind dann identisch, wenn die Paare (x, y) und (x', y') identisch sind, d.h. wenn $x = x'$ und $y = y'$ gelten.

Es sei $z = x + iy \in \mathbb{C}$, dann heißt $\bar{z} = x - iy \in \mathbb{C}$ die zu z *konjugiert komplexe* Zahl. Ist also $z = (x, y)$, so ist $\bar{z} = (x, -y)$.

Das Rechnen mit komplexen Zahlen Für komplexe Zahlen gelten die folgenden Rechenregeln:

1. **Addition:** Es seien $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$ zwei komplexe Zahlen. Dann ist als Summe von z_1 und z_2 die komplexe Zahl

$$z_1 + z_2 = x_1 + x_2 + i(y_1 + y_2). \quad (\text{A.1})$$

erklärt. Die Addition ist also formal der Addition reeller Zahlen äquivalent; sie entspricht der Addition der Vektoren $(x_1, y_1)'$ und $(x_2, y_2)'$.

2. **Multiplikation:** z_1, z_2 seien komplexe Zahlen. Die Multiplikation wird ebenfalls durchgeführt, als multiplizierte man reelle Zahlen: es ist

$$\begin{aligned} z_1 z_2 &= (x_1 + iy_1)(x_2 + iy_2) = x_1 x_2 + ix_1 y_2 + ix_2 y_1 + i^2 y_1 y_2 = \\ &= x_1 x_2 - y_1 y_2 + i(x_1 y_2 + x_2 y_1). \end{aligned} \quad (\text{A.2})$$

Dem Produkt $z = z_1 z_2$ entspricht also das Paar

$$z = (x_1 x_2 - y_1 y_2, x_1 y_2 + x_2 y_1).$$

Es sei insbesondere $z_2 = \bar{z}_1$; dann ist

$$z_1 z_2 = z_1 \bar{z}_1 = |z_1|^2 = x_1^2 + y_1^2 \in \mathbb{R} \quad (\text{A.3})$$

denn $x_1 y_2 + x_2 y_1 = x_1 y_1 - x_1 y_1 = 0$. Der Wert des Betrages der komplexen Zahl $z = x + iy$ folgt hieraus sofort:

$$|z| = \sqrt{x^2 + y^2} \quad (\text{A.4})$$

3. **Division:** Gesucht ist $z = z_1/z_2 = x + iy$. Schreibt man diese Gleichung in der Form $zz_2 = z_1$, so kann man das Produkt zz_2 bilden und daraus die Komponenten des Vektors (x, y) bestimmen. Äquivalent dazu ist das folgende vorgehen: man erweitere z mit $\bar{z}_2$; dann ist $z = z_1\bar{z}_2/(z_2\bar{z}_2) = z_1\bar{z}_2/(x_2^2 + y_2^2)$, nach 2. Setzt man $a = 1/(x_2^2 + y_2^2)$, so ist $z_1/z_2 = az_1\bar{z}_2 = (x_1 + iy_1)(x_2 - iy_2) = x_1x_2 - ix_1y_2 + ix_2y_1 - i^2y_1y_2$, d.h. man hat

$$z = \frac{z_1}{z_2} = x_1x_2 + y_1y_2 - i(x_1y_2 - x_2y_1). \quad (\text{A.5})$$

Die Eulerschen Relationen Nach Abb. A.1 ist

$$\cos \varphi = x/|z|, \quad \sin \varphi = y/|z|$$

und mithin

$$x = |z| \cos \varphi, \quad y = |z| \sin \varphi$$

Für $z = x + iy$ ergibt sich demnach die Darstellung

$$z = |z|(\cos \varphi + i \sin \varphi); \quad (\text{A.6})$$

Diese Darstellung von z heißt *trigonometrische Darstellung* von z , im Gegensatz zur *kartesischen Darstellung* $z = x + iy$. Der Faktor $(\cos \varphi + i \sin \varphi)$ heißt auch *Richtungsfaktor* der komplexen Zahl z . Da zwei komplexe Zahlen $z = (x, y)$ und $z' = (x', y')$ genau dann gleich sind, wenn $x = x'$ und $y = y'$ gilt, folgt, dass dann $|z| = |z'|$ und $\varphi \equiv \varphi' \pmod{2\pi}$ gilt¹

Die Funktionen $\cos \varphi$ und $\sin \varphi$ können in (MacLaurin-) Reihen entwickelt werden. Es ist

$$\cos \varphi = 1 - \frac{\varphi^2}{2!} + \frac{\varphi^4}{4!} - \frac{\varphi^6}{6!} + \dots + (-1)^n \frac{\varphi^{2n}}{2n!} \pm \dots \quad (\text{A.7})$$

$$\sin \varphi = \varphi - \frac{\varphi^3}{3!} + \frac{\varphi^5}{5!} - \frac{\varphi^7}{7!} + \dots + (-1)^n \frac{\varphi^{2n+1}}{(2n+1)!} \pm \dots \quad (\text{A.8})$$

Die Reihenentwicklung für e^λ ist

$$e^\lambda = 1 + \lambda + \frac{\lambda^2}{2!} + \frac{\lambda^3}{3!} + \dots = \sum_{k=0}^{\infty} \frac{\lambda^k}{k!} \quad (\text{A.9})$$

Für $\lambda = i\varphi$ insbesondere erhält man

$$e^{i\varphi} = \sum_{k=0}^{\infty} \frac{(i\varphi)^k}{k!}. \quad (\text{A.10})$$

Berücksichtigt man das eingangs genannte Bildungsprinzip für i^k und schreibt man die Reihen für $\cos \varphi$ und $i \sin(\varphi)$ gemäß (A.7) und (A.8) an, so sieht man, dass $e^{i\varphi}$

¹Es seien a und b zwei reelle Zahlen. $a \equiv b \pmod{2\pi}$, d.h. a ist äquivalent b modulo 2π , bedeutet, dass die Differenz $a - b$ durch 2π ohne Rest teilbar ist. Ist k eine natürliche Zahl, so soll also etwa $(a - b) = k(2\pi)$, also $a = b + k(2\pi)$, gelten.

gerade die Summe der Reihen für $\cos \varphi$ und $i \sin(\varphi)$ ist, d.h. man erhält die *Eulersche Formel*

$$e^{i\varphi} = \cos \varphi + i \sin \varphi. \quad (\text{A.11})$$

Damit erhält man die zu (A.6) äquivalente Schreibweise

$$z = |z|e^{i\varphi}, \quad (\text{A.12})$$

wobei $|z| = \sqrt{x^2 + y^2}$ und $\varphi = \tan^{-1}(y/x)$ ist. Hat man zwei komplexe Zahlen $z_1 = |z_1| \exp(i\varphi_1)$ und $z_2 = |z_2| \exp(i\varphi_2)$, so ergibt sich für das Produkt z sofort

$$z = z_1 z_2 = |z_1| |z_2| e^{i(\varphi_1 + \varphi_2)} = |z| e^{i\varphi} \quad (\text{A.13})$$

Die Multiplikation zweier komplexer Zahlen liefert also eine komplexe Zahl, deren Betrag gleich dem Produkt der Beträge der Faktoren und deren Phasenwinkel φ durch die Summe der Phasenwinkel der Faktoren gegeben ist. Insbesondere erhält man leicht einen Ausdruck für die Potenz z^n , wobei n eine natürliche Zahl ist:

$$z^n = (|z|e^{i\varphi})^n = |z|^n e^{in\varphi} = |z|^n (\cos n\varphi + i \sin n\varphi) \quad (\text{A.14})$$

Der Ausdruck auf der rechten Seite dieser Gleichung ist als *Moivresche Formel* für z^n bekannt.

Für die Division erhält man sofort

$$z = \frac{z_1}{z_2} = \frac{|z_1|}{|z_2|} e^{i(\varphi_1 - \varphi_2)} \quad (\text{A.15})$$

Es sei insbesondere $z_1 = 1$; der Imaginärteil von z_1 sei also gleich Null. Für den Phasenwinkel von z_1 gilt dann $\varphi_1 = \tan^{-1} 0/1 = 0$. Dann folgt aus (A.15)

$$z = \frac{1}{z_2} = \frac{1}{|z_2|} e^{-i\varphi_2} \quad (\text{A.16})$$

Für $z = |z|(\cos \varphi + i \sin \varphi)$ erhält man insbesondere

$$\left| \frac{1}{z} \right| = \frac{1}{|z|}, \quad (\text{A.17})$$

so dass $w = 1/z$ durch

$$w = \frac{1}{z} = \frac{1}{|z|} (\cos \varphi - i \sin \varphi) \quad (\text{A.18})$$

gegeben ist.

Die Eulersche Formel führt zu nützlichen Ausdrücken für $\cos \varphi$ und $\sin \varphi$. Offenbar ist $e^{-i\varphi} = \cos(-\varphi) + i \sin(-\varphi)$. $\cos$ ist eine gerade Funktion und $\sin$ ist eine ungerade Funktion, d.h. $\cos(-\varphi) = \cos \varphi$, $\sin(-\varphi) = -\sin \varphi$. Also ist $e^{-i\varphi} = \cos \varphi - i \sin \varphi$. Demnach findet man $e^{i\varphi} + e^{-i\varphi} = 2 \cos \varphi$, oder

$$\cos \varphi = \frac{e^{i\varphi} + e^{-i\varphi}}{2} \quad (\text{A.19})$$

Ebenso findet man $e^{i\varphi} - e^{-i\varphi} = 2i \sin \varphi$, oder

$$\sin \varphi = \frac{e^{i\varphi} - e^{-i\varphi}}{2i}. \quad (\text{A.20})$$

A.2 Fourier-Reihen

A.2.1 Periodische Funktionen

Definition A.2.1 Es sei $u : \mathbb{R} \rightarrow \mathbb{R}$ eine Funktion mit der Eigenschaft

$$u(x) = u(x) + T, \quad 0 < T \in \mathbb{R} \quad (\text{A.21})$$

für alle x . Dann heißt u periodisch, T ist der Wert der Periode.

Folgerung: Die Funktion u sei periodisch mit der Periode T . Da dann $u(x) = u(x + T)$ für alle $x \in \mathbb{R}$ gilt, kann x durch $x + T$ ersetzt werden, so dass

$$u(x + T) = u((x + T) + T) = u(x + 2T)$$

gelten muß. Ebenso sieht man, dass dann auch $u(x) = u(x + 3T)$ gelten muß, d.h. allgemein

$$u(x) = u(x + nT), \quad n \in \mathbb{N}. \quad (\text{A.22})$$

Man kann sagen, dass periodische Funktionen während verschiedener Zeit- oder Ortsabschnitte der Dauer oder Länge T stets die gleiche Form haben. Repräsentiert x die Zeit, so liegt es nahe, von u als einer Schwingung zu sprechen. Bekannte Beispiele für periodische Funktionen sind $\sin(\omega x)$ und $\cos(\omega x)$. Es werde zunächst der Fall $\omega = 1$ betrachtet. Bekanntlich ist $\sin x = \sin(x + 2\pi)$; ebenso ist $\cos x = \cos(x + 2\pi)$.

Es sei nun $1 \neq \omega \in \mathbb{R}$. Dann existiert sicherlich ein $T \in \mathbb{R}$ derart, dass $\sin(\omega x) = \sin(\omega(x + T))$. Setzt man $\omega x = y$, so gilt sicherlich $\sin y = \sin(y + 2\pi)$; das heißt aber $\sin(\omega x) = \sin(\omega x + 2\pi)$. Also muß $\omega x + 2\pi = \omega(x + T) = \omega x + \omega T$ sein und mithin $\omega T = 2\pi$. Nun existiert sicherlich eine Zahl f derart, dass $\omega = 2\pi f$. Dann folgt $2\pi f T = 2\pi$ oder $f T = 1$, d.h. $f = 1/T$ oder $T = 1/f$. Da $T = 2\pi/\omega = 1/f$ die Dauer einer Periode (Schwingung) ist, muß f die Anzahl der Durchläufe während einer Zeiteinheit sein. f heißt die *Frequenz* der Schwingung. Repräsentiert x einen Ort, so spricht man von f als der *Ortsfrequenz*. Die Betrachtungen für $\cos \omega x$ sind analog.

Es läßt sich zeigen, dass sich beliebige (zumindest alle praktisch relevanten) periodische Funktionen durch *Superposition*, d.h. additive Überlagerung, von Sinus- bzw. Kosinusschwingungen darstellen lassen. Dazu wird zuerst gezeigt, dass die Beziehung

$$A \cos \omega x + B \sin \omega x = |a| \cos(\omega x + \varphi) \quad (\text{A.23})$$

gilt mit

$$|a| = \sqrt{A^2 + B^2}, \quad \varphi = \tan^{-1} \frac{B}{A} \quad (\text{A.24})$$

$|a|$ heißt *Amplitude*, φ heißt die *Phase* der Schwingung. Man zeigt, dass die Beziehung (A.23) tatsächlich gilt, indem man man die Eulerschen Formeln

$$\sin \omega x = (e^{j\omega x} - e^{-j\omega x})/2j, \quad \cos \omega x = (e^{j\omega x} + e^{-j\omega x})/2$$

in die linke Seite von (A.23) einsetzt und die Terme mit den Faktoren e^{jx} bzw. e^{-jx} zusammenfaßt; wegen $B/2j = -jB/2$ erhält man

$$A \cos \omega x + B \sin \omega x = \frac{A + jB}{2} e^{j\omega x} + \frac{A - jB}{2} e^{-j\omega x}$$

Dabei sind $a = A + jB$ und $\bar{a} = A - jB$ konjugiert komplexe Zahlen, die in der Form $a = |a|e^{j\varphi}$, $\bar{a} = |a|e^{-j\varphi}$ mit $|a| = \sqrt{A^2 + B^2}$, $\varphi = \tan^{-1}(B/A)$ dargestellt werden können. Setzt man diese Beziehungen ein, erhält man (A.23). $\square$

Die Funktion $u(x) = |a| \cos(\omega x + \varphi)$ ist wieder periodisch, mit der gleichen Periode $T = \omega/2\pi$ wie $\cos \omega x$. Die Tatsache, dass u um den Betrag φ (phasen-)versetzt ist, ändert ja nichts an der Periodizität.

Es werden nun zwei Schwingungen $u_1 = \cos \omega_1 x$ und $u_2 = \cos \omega_2 x$ betrachtet, für die $\omega_2 = 2\omega_1$ gilt. Dann folgt $2\pi f_2 = 2\pi f_1$, d.h. u_2 hat die doppelte Frequenz wie u_1 . u_1 hat die Periode $T_1 = \omega_1/2\pi = 1/f_1$, und u_2 hat die Periode $T_2 = \omega_2/2\pi = 1/f_2 = T_1/2$. Natürlich hat u_2 auch die Periode T_1 , denn nach der Zeit T_1 wiederholt sich der Schwingungsverlauf ja ebenfalls wieder. Analog findet man, dass $u_3 = \cos \omega_3 x$ mit $\omega_3 = 3\omega_1$ die dreifache Frequenz $f_3 = 3f_1$ und dementsprechend die Periode $T_3 = T_1/3$ hat. Auch u_3 hat gleichzeitig die Periode T_1 , da sich u_3 ja ebenfalls nach der Zeit T_1 (zum dritten Mal!) wiederholt. Allgemein hat $\cos \omega_n x = \cos n\omega_1 x$ die Periode $T_n = T_1/n$ und die Frequenz $f_n = nf_1$ und wiederholt sich nach der Zeit T_1 zum n -ten Mal, ist damit also auch wieder periodisch mit der Periode T_1 .

Gelten nun für die Funktionen $u_1 = \cos \omega_1 x$ und $u_2 = \cos \omega_2 x$ die Beziehungen $\omega_1 = m\omega_0$, $\omega_2 = n\omega_0$, so haben beide Schwingungen u.a. die Periode $T_0 = 2\pi/\omega_0$, da sich dann nach der Zeit T_0 der Verlauf beider Funktionen wiederholt. $\omega_0 = 2\pi f_0$ heißt die *Grundschwingung*, bzw. f_0 heißt die *Grundfrequenz* von u_1 bzw. u_2 , und wegen $m\omega_0 = \omega_m = 2\pi m f_0$, $n\omega_0 = \omega_n = 2\pi n f_0$ heißen ω_m und ω_n für $m > 1$ die $(m-1)$ -ten bzw. $(n-1)$ -ten *Oberschwingungen* oder *Harmonischen Oberschwingungen* von ω_0 . f_m und f_n sind die entsprechenden harmonischen Frequenzen.

Definition A.2.2 Für die (Kreis-)Frequenzen ω_1 und ω_2 gelte $\omega_1 = m\omega_0$ und $\omega_2 = n\omega_0$, so dass

$$\frac{\omega_1}{\omega_2} = \frac{m}{n} \quad (\text{A.25})$$

gilt. Dann heißen ω_1 und ω_2 *kommensurabel*.

Sind zwei Schwingungen mit den Frequenzen ω_1 und ω_2 *nicht* kommensurabel, so existiert keine Periode T , nach der für beide Schwingungen der Verlauf von vorn beginnt. Sind sie dagegen kommensurabel, so existiert diese Perioden T ; für die eine Schwingung beginnt nach der Zeit T der Verlauf etwa zum c -ten, für die andere zum n -ten Mal, $1 \leq m, n \in \mathbb{N}$.

A.2.2 Superpositionen

Es sei $A = a_n$, $B = b_n$, $\omega = n\omega_0$; dann folgt aus (A.23) die Beziehung

$$u_n(x) = a_n \cos n\omega_0 x + b_n \sin n\omega_0 x = |c_n| \cos(n\omega_0 + \varphi_n) \quad (\text{A.26})$$

mit

$$|c_n| = \sqrt{a_n^2 + b_n^2}, \quad \varphi_n = \tan^{-1} \frac{b_n}{a_n} \quad (\text{A.27})$$

Zum Zeitpunkt $T_0 = \omega_0/2\pi$ beginnt für alle u_n , $n = 1, 2, 3, \dots$ ein erneuter Durchlauf. Deshalb wird auch die Summe

$$S(x) = \alpha + \sum_{k=1}^n (a_k \cos k\omega_0 x + b_k \sin k\omega_0 x) = \alpha + \sum_{k=1}^n |c_k| \cos(k\omega_0 x - \varphi_k) \quad (\text{A.28})$$

ein periodische Funktion sein, die nach der Zeit T_0 ihren Durchlauf erneuert. Die additive Konstante α ändert an dieser Periodizität nichts, da α als eine Funktion aufgefaßt werden kann, die an jeder Stelle periodisch ist.

Aufgrund der Eulerschen Relationen gilt

$$\cos k\omega_0 x = \frac{e^{j(k\omega_0 x)} + e^{-j(k\omega_0 x)}}{2}, \quad \sin k\omega_0 x = \frac{e^{jk\omega_0 x} - e^{-jk\omega_0 x}}{2j}$$

Dann folgt

$$a_k \cos k\omega_0 x + b_k \sin k\omega_0 x = \left(\frac{a_k}{2} + \frac{b_k}{2j} \right) e^{jk\omega_0 x} + \left(\frac{a_k}{2} - \frac{b_k}{2j} \right) e^{-jk\omega_0 x}.$$

Die Faktoren $a_k/2 + b_k/2j$ und $a_k/2 - b_k/2j$ sind offenbar komplexe, insbesondere konjugiert komplexe Zahlen. Führt man die Notation

$$\alpha_k := \frac{a_k}{2} + \frac{b_k}{2j} = \frac{a_k}{2} - \frac{jb_k}{2} = \frac{1}{2}(a_k - jb_k) \quad (\text{A.29})$$

$$\alpha_{-k} := \frac{a_k}{2} - \frac{b_k}{2j} = \frac{a_k}{2} + \frac{jb_k}{2} = \frac{1}{2}(a_k + jb_k) \quad (\text{A.30})$$

ein, so lassen sich die Koeffizienten a_k und b_k in der Form

$$a_k = \alpha_k + \alpha_{-k}, \quad b_k = j(\alpha_k - \alpha_{-k}) \quad (\text{A.31})$$

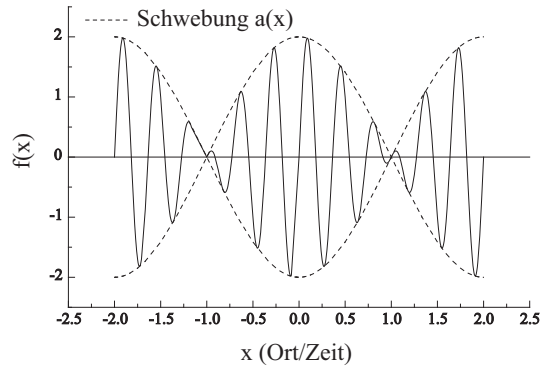
anschreiben, und die in (A.28) eingeführte Summe $S(x)$ kann in der Form

$$S(x) = \sum_{k=-n}^n \alpha_k e^{jk\omega_0 x} \quad (\text{A.32})$$

geschrieben werden.

Der Verlauf der Funktion $S(x)$ wird von den Konstanten $\omega_0 = 2\pi f_0$ und a_k, b_k , $k = 1, 2, \dots, n$, also von insgesamt $2n + 1$ Konstanten bestimmt. Diese Konstanten kann man (nahezu) beliebig wählen, um dann eine entsprechend willkürliche Funktion $S(x)$ zu erhalten. Es liegt deshalb nahe, zu vermuten, dass eine vorgegebene Funktion $u(x)$ durch geeignete Wahl dieser Konstanten ebenfalls als eine solche Reihe dargestellt werden kann. Um diese Vermutung diskutieren zu können, werden in Abschnitt A.2.3 einige Eigenschaften von Sin- und Kosinusfunktionen angegeben.

Abbildung A.2: Schwebung, $\omega_1 = 3$, $\omega_2 = 2.5$



Anmerkung: Schwebungen (*Beats*) Es sind Summen der Form $a_k \cos k\omega_0 x + b_k \sin k\omega_0 x$ betrachtet worden. Bei Summen der Form

$$y(x) = \sin \omega_1 x + \sin \omega_2 x \quad (\text{A.33})$$

ergeben sich neue Phänomene. Anhand der Eulerschen Formeln rechnet man leicht nach, dass dieser Ausdruck gleichbedeutend mit

$$y(x) = 2 \cos \left(\frac{\omega_1 - \omega_2}{2} x \right) \sin \left(\frac{\omega_1 + \omega_2}{2} x \right) = a(x) \sin \left(\frac{\omega_1 + \omega_2}{2} x \right) \quad (\text{A.34})$$

ist, wobei

$$a(x) = 2 \cos \left(\frac{\omega_1 - \omega_2}{2} x \right)$$

ist. Dementsprechend ist $y(x)$ eine Sinusschwingung mit der Frequenz $(\omega_1 + \omega_2)/2$, d.h. mit der Periode $T = 2\pi/((\omega_1 + \omega_2)/2) = 4\pi/(\omega_1 + \omega_2)$, mit der *von x abhängenden* Amplitude $a(x) = 2 \cos((\omega_1 - \omega_2)x/2)$. Die Amplitude verändert sich also wie eine Kosinusfunktion mit der Frequenz $(\omega_1 - \omega_2)/2$. Die Periode dieser Schwingung ist $4\pi/(\omega_1 - \omega_2)$. Sind ω_1 und ω_2 beide "groß", liegen aber nahe beieinander, so ist die Differenz $\omega_1 - \omega_2$ aber "klein", so ist die Periode der Amplitudenschwingung groß im Vergleich zur Schwingung mit der Frequenz $(\omega_1 + \omega_2)/2$. Die Amplitudenschwingung heißt *Schwebung* (vergl. Abb. A.2).

A.2.3 Die Orthogonalitätsrelationen der Sinus- und Kosinusfunktionen

Definition A.2.3 Es seien $\phi_1, \phi_2, \dots$ Funktionen, für die

$$\int_a^b \phi_m(x) \phi_n(x) dx = \begin{cases} 0, & m \neq n \\ r_n, & m = n \end{cases} \quad (\text{A.35})$$

gilt. Dann heißen die Funktionen orthogonal auf dem Intervall (a, b) .

Es sei $T = 2\pi/\omega_0$. Insbesondere für Sinus- und Kosinusfunktionen gelten dann die folgenden Relationen

$$\int_{-T/2}^{T/2} \cos m\omega_0 x dx = 0, \quad m \neq 0 \quad (\text{A.36})$$

$$\int_{-T/2}^{T/2} \sin m\omega_0 x dx = 0, \quad m \neq 0 \quad (\text{A.37})$$

$$\int_{-T/2}^{T/2} \cos m\omega_0 x \cos n\omega_0 x dx = \begin{cases} 0, & m \neq n \\ T/2, & m = n \neq 0 \end{cases} \quad (\text{A.38})$$

$$\int_{-T/2}^{T/2} \sin m\omega_0 x \sin n\omega_0 x dx = \begin{cases} 0, & m \neq n \\ T/2, & m = n \neq 0 \end{cases} \quad (\text{A.39})$$

$$\int_{-T/2}^{T/2} \sin m\omega_0 x \cos n\omega_0 x dx = 0, \quad \text{für alle } m, n \in \mathbb{N} \quad (\text{A.40})$$

Diese Relationen, insbesondere die Relationen (A.38) - (A.40), heißen die *Orthogonalitätsrelationen*; die Sinus- und Kosinusfunktionen sind für unterschiedliche Frequenzparameter offenbar orthogonal. Man prüft dies nach, indem man von den Eulerschen Relationen Gebrauch macht; die Rechnungen sollen hier nicht im Einzelnen durchgeführt werden.

Fourierreihenentwicklung für eine willkürliche periodische Funktion.

In (A.28) ist die Summe

$$S(x) = \alpha + \sum_{k=1}^n (a_k \cos k\omega_0 x + b_k \sin k\omega_0 x)$$

eingeführt wurde. Die Orthogonalitätsrelationen liefern dann

$$\frac{1}{\pi} \int_{-\pi}^{\pi} S(x) \cos kx dx = a_k, \quad k = 1, 2, \dots \quad (\text{A.41})$$

$$\frac{1}{\pi} \int_{-\pi}^{\pi} S(x) \sin kx dx = b_k, \quad k = 1, 2, \dots \quad (\text{A.42})$$

d.h. man erhält die Koeffizienten a_k und b_k , indem man $S(x)$ mit $\cos kx$ bzw. $\sin kx$ multipliziert und über das Intervall $(-\pi, \pi)$ integriert.

Die Parameter α , a_k und b_k sind willkürlich gewählt worden; es ergibt sich dementsprechend eine Funktion $S(x)$. Gemäß (A.41) und (A.42) kann man aber aus $S(x)$ die Koeffizienten a_k und b_k wiederum errechnen, indem $S(x)$ mit $\cos kx$ bzw. $\sin kx$ multipliziert und das Produkt $S(x) \cos kx$ und $S(x) \sin kx$ über dem Intervall $(-\pi, \pi)$ integriert. Dies legt nahe, dass man für eine beliebige, vorgegebene Funktion $u(x)$ Koeffizienten a_k und b_k errechnen kann derart, dass sich $u(x)$ durch eine Reihe

$$u(x) = \alpha + \sum_{k=1}^{\infty} (a_k \cos k\omega_0 x + b_k \sin k\omega_0 x) \quad (\text{A.43})$$

darstellen läßt. Dabei soll also

$$a_k = \frac{1}{\pi} \int_{-\pi}^{\pi} u(x) \cos k\omega_0 x \, dx, \quad b_k = \frac{1}{\pi} \int_{-\pi}^{\pi} u(x) \sin k\omega_0 x \, dx \quad (\text{A.44})$$

und

$$\alpha = \frac{a_0}{2} = \frac{1}{2\pi} \int_{-\pi}^{\pi} S(x) \, dx \quad (\text{A.45})$$

gelten. Es sei darauf hingewiesen, dass hier die Summation bis unendlich geht; es ist ja a priori nicht bekannt, wieviele Summanden benötigt werden.

Man muß nun beweisen, dass $u(x)$ tatsächlich in der Form (A.43) darstellbar ist. Ein korrekter Beweis ist ein wenig länglich und soll deshalb hier nicht dargestellt werden, man findet ihn etwa bei Courant (1961)².

Beispiel A.2.1 Es sei $u(t) = 1$ für $0 < t \leq T/2$, $u(t) = -1$ für $-T/2 < t \leq 0$ und $T/2 < t \leq T$, etc (vergl. Fig. 6.1). Man findet

$$\frac{1}{2}a_0 = \frac{1}{T} \int_{-T/2}^{T/2} f(t) \, dt = 0, \quad (\text{A.46})$$

$$a_n = \frac{2}{T} \int_{-T/2}^{T/2} u(t) \cos(n\omega t) \, dt = 0, \quad \text{für alle } n > 1, \text{ und} \quad (\text{A.47})$$

$$b_n = \frac{2}{T} \int_{-T/2}^{T/2} u(t) \sin(n\omega t) \, dt = \begin{cases} 0, & n \text{ gerade,} \\ 2(1 - \cos n\pi)/n\pi, & n \text{ ungerade} \end{cases} \quad (\text{A.48})$$

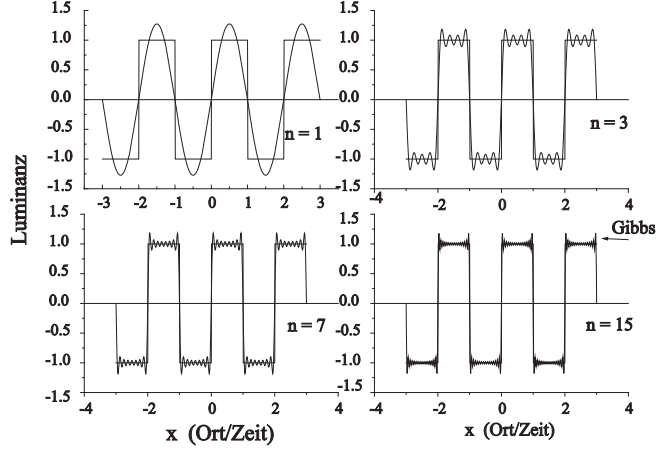
Damit ist

$$u(x) = \frac{4}{\pi} \left(\sin \omega x + \frac{1}{3} \sin 3\omega x + \frac{1}{5} \sin 5\omega x + \dots \right), \quad (\text{A.49})$$

das Rechteckgitter setzt sich also aus den ungeraden Harmonischen zusammen. Abb. A.3 zeigt das Rechteckgitter und seine Approximation durch die Superposition von Sinusschwingungen für verschiedene Werte von n .

²Courant, R. Vorlesungen über Differential- und Integralrechnung. Erster Band: Funktionen einer Veränderlichen. Springer-Verlag Berlin, Göttingen, Heidelberg 1961

Abbildung A.3: Rechteckgitter



Es sei andererseits $u(x) = 1$ für $-T/4 \leq x \leq T/4$, $u(x) = -1$ für $-3T/4 \leq x \leq -T/4$ und $T/4 \leq x \leq 3T/4$, etc. Dieses Rechteckgitter ist eine *gerade* Funktion, d.h. es gilt $u(-x) = u(x)$ für alle x . Man findet die Darstellung

$$u(x) = \frac{4}{\pi} \left(\cos \omega_0 x - \frac{1}{3} \cos 3\omega_0 x + \frac{1}{5} \cos 5\omega_0 x - \dots \right) \quad (\text{A.50})$$

mit $\omega_0 = 2\pi/T$. □

A.3 Fouriertransformierte

A.3.1 Definition der Fouriertransformierten

Gegeben sei nun eine nichtperiodische Funktion $f(t)$. Um ihre Zerlegung in harmonische Funktionen zu erhalten, werde eine periodische Funktion F betrachtet mit der Eigenschaft, dass $F(t) = f(t)$ für $-T/2 < t < T/2$ gilt. F kann in eine Fourierreihe gemäß (A.28) entwickelt werden. Für $T \rightarrow \infty$ erhält man daraus eine Fourier-Repräsentation für $F(t)$. In der Tat, gemäß (A.31) und (A.44) man

$$F(t) = \sum_{n=-\infty}^{\infty} c_n e^{in\omega t}, \quad c_n = \frac{1}{T} \int_{-T/2}^{T/2} f_T(t) e^{-in\omega t}, \quad \omega = 2\pi/T. \quad (\text{A.51})$$

Substituiert man das Integral für c_n in die Summe, so resultiert

$$f_T(t) = \sum_{n=-\infty}^{\infty} \left[\frac{1}{2\pi} \int_{-T/2}^{T/2} f_T(\tau) e^{-in\omega\tau} d\tau \right] \omega e^{in\omega t},$$

wobei $1/T$ durch $\omega/2\pi$ ersetzt wurde. Aus der Definition von ω folgt $\omega \rightarrow 0$ für $T \rightarrow \infty$. Es sei $\omega_0 = \Delta\omega$; dann ist $n\Delta\omega = n\omega_0$. Es werde nun angenommen, dass

Abbildung A.4: Pulsfolgen mit verschiedener Periodizität

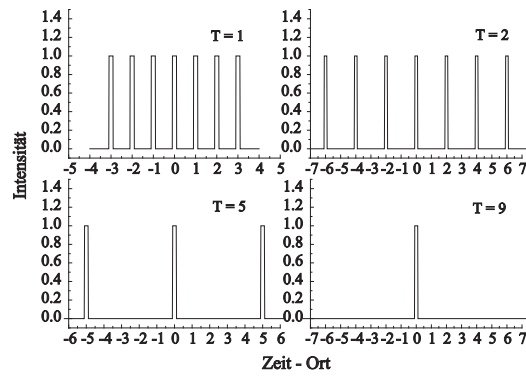
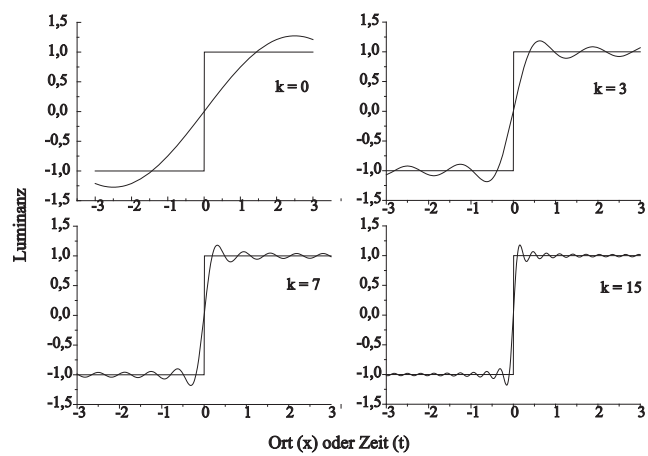


Abbildung A.5: Approximation einer Schrittfunktion



für $\omega_0 \rightarrow 0$ und $n \rightarrow \infty$ das Produkt $n\omega_0 = n\Delta\omega$ endlich bleibt: $n\Delta\omega \rightarrow \omega$. Die Darstellung für f_T wird dann

$$f_T(t) = \sum_{n=-\infty}^{\infty} \left[\frac{1}{2\pi} \int_{-T/2}^{T/2} f_T(\tau) e^{-i\Delta\omega\tau} d\tau \right] e^{in\Delta\omega t} \Delta\omega,$$

und für $T \rightarrow \infty$, $\Delta\omega \rightarrow d\omega$ geht die Summe in ein Integral über:

$$f_T(t) = f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(\tau) e^{-i\omega\tau} e^{i\omega t} d\tau. \quad (\text{A.52})$$

Man definiert nun

$$F(\omega) = \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt, \quad (\text{A.53})$$

und aus (A.52) wird

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega t} d\omega. \quad (\text{A.54})$$

Der Ausdruck (A.54) ist die Fourier-Darstellung für eine nichtperiodische Funktion $f(t)$; $F(\omega)$ ist die Fouriertransformierte von $f(t)$.

$F(\omega)$ ist im allgemeinen eine komplexe Funktion, d.h. $F(\omega) = F_1(\omega) + iF_2(\omega)$; also kann F in der Form

$$F(\omega) = |F(\omega)| e^{i\phi} \quad (\text{A.55})$$

dargestellt werden, mit

$$|F(\omega)| = \sqrt{F_1^2(\omega) + F_2^2(\omega)}, \quad \phi = \tan^{-1}(F_2/F_1). \quad (\text{A.56})$$

$F(\omega)$ existiert nicht für alle Funktionen $f(t)$; eine hinreichende Bedingung für die Existenz einer Fouriertransformierten ist

$$\int_{-\infty}^{\infty} |f(t)| dt < \infty, \quad (\text{A.57})$$

d.h. f muß absolut integrierbar sein.

A.3.2 Einige allgemeine Sätze

1. Kosinus- und Sinus-Transformierte
2. Symmetrie
3. Skalierung
4. Modulierung
5. Parsevals Satz

Satz A.3.1 *Es sei*

$$R(\omega) = \int_{-\infty}^{\infty} f(x) \cos \omega x \, dx \quad (\text{A.58})$$

$$X(\omega) = \int_{-\infty}^{\infty} f(x) \sin \omega x \, dx \quad (\text{A.59})$$

Dann kann F in der Form

$$F(\omega) = R(\omega) - iX(\omega). \quad (\text{A.60})$$

geschrieben werden.

Anmerkung: $R(\omega)$ heißt *Kosinus-Transformierte* und $X(\omega)$ heißt *Sinus-Transformierte* der reellen Funktion f .

Beweis: Aufgrund der Definition von $e^{-i\omega x}$ kann (A.53) in der Form

$$F(\omega) = \int_{-\infty}^{\infty} f(x) e^{-i\omega x} \, dx = \int_{-\infty}^{\infty} f(x) \cos \omega x \, dx - i \int_{-\infty}^{\infty} f(x) \sin \omega x \, dx \quad (\text{A.61})$$

geschrieben werden. □

Es ist

$$\cos(-\omega x) = \cos(\omega x),$$

d.h. der Kosinus ist eine *gerade Funktion*. Andererseits ist

$$\sin(-\omega x) = -\sin(\omega x),$$

so dass der Sinus eine *ungerade Funktion* ist. Hieraus folgt sofort

$$R(-\omega) = R(\omega), \quad X(-\omega) = -X(\omega) \quad (\text{A.62})$$

Die Fouriertransformierte F ist i.a. eine komplexe Funktion, so dass $F(\omega) = F_1(\omega) + iF_2(\omega)$, wobei F_1 der Realteil und F_2 der Imaginärteil der Funktion ist. $F^* = F_1 - iF_2$ ist dann die zu F konjugiert komplexe Funktion. Aus (A.62) folgt dann

$$F^*(\omega) = F(-\omega) \quad (\text{A.63})$$

Es sei insbesondere f eine gerade Funktion, so dass $f(-x) = f(x)$. Man sieht leicht, dass dann $X(\omega) = 0$ ist, denn

$$X(\omega) = \int_{-\infty}^{\infty} f(x) \sin \omega x \, dx = \int_{-\infty}^0 f(x) \sin \omega x \, dx + \int_0^{\infty} f(x) \sin \omega x \, dx = 0,$$

denn

$$\int_{-\infty}^0 f(x) \sin \omega x \, dx = - \int_0^{\infty} f(x) \sin \omega x \, dx$$

Also ist die Fouriertransformierte einer reellen, geraden Funktion durch

$$F(\omega) = \int_{-\infty}^{\infty} f(x) \cos \omega x \, dx = 2 \int_0^{\infty} f(x) \cos \omega x \, dx \quad (\text{A.64})$$

gegeben. Umgekehrt folgt, dass $R(\omega) = 0$, wenn f ungerade ist, und es gilt dann

$$F(\omega) = \int_{-\infty}^{\infty} f(x) \sin \omega x \, dx = -2 \int_0^{\infty} f(x) \sin \omega x \, dx \quad (\text{A.65})$$

Satz A.3.2 (*Symmetrie*) Die Funktion $f(x)$ habe die Fouriertransformierte $F(\omega)$. Dann hat die Funktion $F(x)$ die Fouriertransformierte $2\pi f(-\omega)$.

Beweis: Es ist

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega x} \, d\omega$$

so dass

$$2\pi f(x) = \int_{-\infty}^{\infty} F(\omega) e^{i\omega x} \, d\omega$$

Geht man von x zu $-x$ über, so erhält man

$$2\pi f(-x) = \int_{-\infty}^{\infty} F(\omega) e^{-i\omega x} \, d\omega$$

Aber sowohl x wie ω sind reelle Zahlen; vertauscht man die Bezeichnung dieser Zahlen, so erhält man

$$2\pi f(-\omega) = \int_{-\infty}^{\infty} F(x) e^{-ix\omega} \, dx$$

und man sieht, dass die rechte Seite gerade die Fouriertransformierte der Funktion $F(x)$ darstellt. $\square$

Satz A.3.3 (*Skalierung*) Es sei $f(x)$ eine Funktion mit der Fouriertransformierten $F(\omega)$. Dann gilt

$$F(f(ax)) = \frac{1}{|a|} F\left(\frac{\omega}{a}\right) \quad (\text{A.66})$$

Beweis: Es sei $y = ax$ und damit $x = y/a$; dann ist $dy/dx = a$ und demnach $dx = dy/a$. Mithin ist für $a > 0$

$$\begin{aligned} F(f(y)) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} f(y) e^{-i\omega y} \, dy \\ &= \frac{1}{2\pi} \frac{1}{a} \int_{-\infty}^{\infty} f\left(\frac{x}{a}\right) \exp\left(-i\omega \frac{x}{a}\right) \, dx \\ &= \frac{1}{2\pi} \frac{1}{a} \int_{-\infty}^{\infty} f\left(\frac{x}{a}\right) \exp\left(-i\frac{\omega}{a} x\right) \, dx \\ &= \frac{1}{a} F\left(\frac{\omega}{a}\right) \end{aligned} \quad (\text{A.67})$$

Für $a < 0$ kommt man zum gleichen Ergebnis, und damit ist die Aussage bewiesen. $\square$

Satz A.3.4 (*Shift-Theorem*) Es sei $f(x)$ eine Funktion mit der Fouriertransformierten $F(\omega)$. Dann folgt

$$F(f(x - x_0)) = F(\omega)e^{-i\omega x_0} \quad (\text{A.68})$$

Beweis: Es sei $y = x - x_0$. Dann ist $dy/dx = 1$, d.h. $dy = dx$ und

$$\begin{aligned} F(f(x - x_0)) &= \int_{-\infty}^{\infty} f(x - x_0)e^{-i\omega x} dx \\ &= \int_{-\infty}^{\infty} f(y)e^{-i\omega(y+x_0)} dy \\ &= e^{-i\omega x_0} \int_{-\infty}^{\infty} f(y)e^{-i\omega y} dy \\ &= e^{-i\omega x_0} F(\omega) \end{aligned}$$

Dem Satz A.3.4 entspricht der

Satz A.3.5 (*Frequency-Shift-Theorem*) Es sei $g(x) = f(x)e^{i\omega_0 x}$ und f habe die Fouriertransformierte $F(\omega)$. Dann gilt

$$F(g(x)) = F(f(x)e^{i\omega_0 x}) = F(\omega - \omega_0) \quad (\text{A.69})$$

Beweis: Es ist

$$\begin{aligned} F(f(x)e^{i\omega_0 x}) &= \int_{-\infty}^{\infty} (f(x)e^{i\omega_0 x}) e^{-i\omega x} dx \\ &= \int_{-\infty}^{\infty} f(x)e^{-i(\omega - \omega_0)x} dx \end{aligned}$$

$\square$

Satz A.3.6 (*Modulierung*) Die Funktion f habe die Fouriertransformierte F . Dann gilt

$$F(f(x) \cos \omega_0 x) = \frac{1}{2}(F(\omega + \omega_0) + F(\omega - \omega_0)) \quad (\text{A.70})$$

Beweis: Es ist $\cos \omega_0 x = (e^{i\omega_0 x} + e^{-i\omega_0 x})/2$; (A.70) folgt dann aus dem Frequenz-Shift-Theorem, d.h. aus Satz A.3.4. $\square$

Satz A.3.7 Es sei f eine mindestens n -mal differenzierbare Funktion von x und F sei die zugehörige Fouriertransformierte. Dann gilt

$$F((-ix)^n f(x)) = F^{(n)}(\omega), \quad f^{(n)}(x) = (i\omega)^n F(\omega) \quad (\text{A.71})$$

Beweis: Die Aussage folgt sofort, wenn man F in (A.53) bzw. f in (A.54) n -mal differenziert. $\square$

Satz A.3.8 *Es seien f und g zwei Funktionen von x mit den Fouriertransformierten F und G . Mit $f * g$ werde die Faltung von f und g bezeichnet, d.h. es sei*

$$f * g = \int_{-\infty}^{\infty} f(\xi)g(x - \xi) d\xi$$

Dann gilt

$$F(f * g) = F(\omega)G(\omega) \quad (\text{A.72})$$

*Ist andererseits $F * G$ die Faltung von F und G , so gilt*

$$F(fg) = \frac{1}{2\pi} F(\omega) * G(\omega) \quad (\text{A.73})$$

Anmerkung: Die Aussage (A.73) heißt auch *Faltung im Frequenzbereich*.

Beweis: Es ist

$$F(f * g) = \int_{-\infty}^{\infty} e^{-i\omega x} \left(\int_{-\infty}^{\infty} f(\xi)g(x - \xi) d\xi \right) dx$$

Setzt man $x = \xi + y$ und vertauscht die Reihenfolge der Integration, so folgt

$$\int_{-\infty}^{\infty} f(\xi) \int_{-\infty}^{\infty} e^{-i\omega(\xi+y)} g(y) dy d\xi = \int_{-\infty}^{\infty} f(\xi) e^{-i\omega\xi} d\xi \int_{-\infty}^{\infty} g(y) e^{-i\omega y} dy,$$

und dies ist die Behauptung. (A.73) folgt auf analoge Weise. $\square$

Satz A.3.9 (*Parsevals Satz*) *Es seien f und g beliebige, d.h. nicht notwendig reelle Funktionen und F, G seien die entsprechenden Fouriertransformierten. Dann gilt die Beziehung*

$$\int_{-\infty}^{\infty} f(x)g^*(x) dx = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega)G^*(\omega) d\omega \quad (\text{A.74})$$

Beweis: Man setzt in (A.73) $x = 0$ und die Aussage (A.74) folgt. $\square$

Satz A.3.10 *Die Funktion f habe die Fouriertransformierte F . Dann gilt*

$$\int_{-\infty}^{\infty} |f(x)|^2 dx = \frac{1}{2\pi} \int_{-\infty}^{\infty} |F(\omega)|^2 d\omega \quad (\text{A.75})$$

Beweis: Die Aussage folgt sofort aus (A.74) für den Spezialfall $f(x) \equiv g(x)$. $\square$

Anmerkung: das Integral $\int |f(x)|^2 dx$ bezeichnet die Energie des Signals f . Nach (A.3.10) ist die Energie als Integral über dem Quadrat der Amplituden der Frequenzen gegeben, durch die f definiert wird.

A.3.3 Die Dirac-Delta Funktion

Definition A.3.1 Die Funktion $\delta(x)$ sei durch die Eigenschaften

$$\delta(x) = \begin{cases} 0, & x \neq 0 \\ \infty, & x = 0 \end{cases}, \quad \text{und} \quad \int_{-\infty}^{\infty} \delta(x) dx = 1 \quad (\text{A.76})$$

definiert. Dann heit $\delta(x)$ die Dirac-Delta-Funktion oder auch kurz die δ -Funktion.

Die δ -Funktion ist eine *verallgemeinerte Funktion*; ihre Eigenschaften sind ja ein wenig bizarr. Sie kann als Grenzfall einiger bekannter Funktionen gesehen werden. Bekanntlich hat zum Beispiel das Integral der Gauschen Dichte

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{x^2}{2\sigma^2}\right)$$

den Wert 1 fr alle $\sigma > 0$. Fr $\sigma \rightarrow 0$ geht $f(x)$ dann in die δ -Funktion ber. Ein anderes Beispiel ist die Funktion

$$f(x) = \begin{cases} 1/\Delta x, & -\Delta x/2 \leq x \leq \Delta x/2 \\ 0, & x < -\Delta x/2 \quad \text{oder} \quad x > \Delta x/2 \end{cases} \quad (\text{A.77})$$

Offenbar ist

$$\int_{-\infty}^{\infty} f(x) dx = \Delta x \frac{1}{\Delta x} = 1$$

fr alle $\Delta x > 0$, und fr $\Delta x \rightarrow 0$ geht $f(x)$ in die δ -Funktion ber. Es gilt

$$\int_{-\infty}^{\infty} f(x) \delta(x - x_0) dx = f(x_0) \quad (\text{A.78})$$

Fr die Fouriertransformierte der δ -Funktion folgt deshalb

$$F(\delta(x)) = \int_{-\infty}^{\infty} \delta(x) e^{-i\omega x} dx = e^{-i\omega x}|_{x=0} = 1 \quad (\text{A.79})$$

Diese Beziehung ist eine direkte Anwendung von (A.78), das Integral ist gleich $\exp(-i\omega x)$ an der Stelle $x = 0$, also gleich 1. (A.79) bedeutet, dass die zeitlich oder rtlich nicht ausgedehnte Funktion δ alle Frequenzen mit der gleichen Amplitude 1 enthlt. Eine unmittelbare Folgerung ist

$$\delta(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} 1 \cdot e^{i\omega x} d\omega = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{i\omega x} d\omega \quad (\text{A.80})$$

Das erste Integral auf der rechten Seite ist einfach der Ausdruck fr die inverse Fouriertransformierte. Das zweite Integral ist eine Aussage ber $\exp(i\omega x)$. Es ist klar, dass dieses Integral nicht im blichen Sinne hergeleitet werden kann, denn die Bedingung $\int |f(x)| dx < \infty$ (vergl. (A.57)) ist fr $f(x) = \exp(i\omega x)$ nicht erfllt. Dieser Sachverhalt ist eine Konsequenz der Definition der Dirac-Delta-Funktion. Es zeigt sich aber, dass *formale* Rechnungen mit (A.80) i.a. zu korrekten Ergebnissen

führen, wie etwa die weiter unten eingeführten Fouriertransformierten der sin- und cos-Funktionen zeigen.

Aus dem Shift-Theorem folgt sofort

$$F(\delta(x - x_0)) = \int_{-\infty}^{\infty} \delta(x - x_0) e^{-i\omega x} dx = e^{-i\omega x}|_{x=x_0} = e^{-i\omega x_0} \quad (\text{A.81})$$

Aus der Bedingung (A.57) folgt, dass die Funktionen $\sin \omega_0 x$ und $\cos \omega_0 x$ keine Fouriertransformierte im Sinne einer "gewöhnlichen" Funktion haben können. Andererseits sollte z.B. $\sin \omega_0 x$ aus Sinusfunktionen zusammensetzbar sein, da diese Funktion ja selbst eine solche Funktion ist.

A.4 Fourier-Transformationen spezieller Funktionen

A.4.1 Die Sinus- bzw. Kosinusfunktion

Es sei $s(t) = \sin(\omega t + \phi)$; gesucht ist die Fouriertransformierte $S(j\omega)$ von s . Es ist

$$\sin(\omega_0 t + \phi) = \frac{1}{2j} \left(e^{j\omega_0 t + \phi} - e^{-j\omega_0 t - \phi} \right), \quad j = \sqrt{-1}. \quad (\text{A.82})$$

Dann ist

$$\begin{aligned} S(j\omega) &= \int_{-\infty}^{\infty} \sin(\omega_0 t + \phi) e^{-j\omega t} dt = \frac{1}{2j} \left[\int_{-\infty}^{\infty} e^{j(\omega_0 - \omega)t + j\phi} dt - \int_{-\infty}^{\infty} e^{-j(\omega_0 - \omega)t - j\phi} dt \right] \\ &= \frac{2\pi e^{j\phi}}{2j} \delta(\omega_0 - \omega) - \frac{2\pi e^{-j\phi}}{2j} \delta(\omega_0 + \omega). \end{aligned} \quad (\text{A.83})$$

$$= \quad (\text{A.84})$$

A.4.2 Die Exponentialfunktion

Es sei $f(t) = ae^{-\alpha t}$, $\alpha > 0$. Es ist $\int |ae^{-\alpha t}| dt = a(1 - e^{-\alpha t})/\alpha \rightarrow a/\alpha$ für $t \rightarrow \infty$, also ist f absolut integrierbar und die Fouriertransformierte existiert. Nach (A.53) ist dann

$$F(\omega) = a \int_0^{\infty} e^{-\alpha t - i\omega t} dt = a \int_0^{\infty} e^{-(\alpha + i\omega)t} dt = \frac{a}{\alpha + i\omega}. \quad (\text{A.85})$$

Es werde noch die n -fache Faltung der Exponentialfunktion $f(t) = a \exp(-\alpha t)$, $\alpha > 0$, $a > 0$ betrachtet; sie entspricht n rückwirkungsfrei hintereinandergeschalteten Systemen mit identischer Impulsantwort $f(t)$. Man findet

$$f_n(t) = \frac{a^n}{n!} t^{n-1} e^{-\alpha t}. \quad (\text{A.86})$$

Die Fouriertransformierte von f_n ergibt sich

$$F_n(\omega) = \left(\frac{a}{\alpha + i\omega} \right)^n. \quad (\text{A.87})$$

Abbildung A.6: Faltungen der Exponentialfunktion

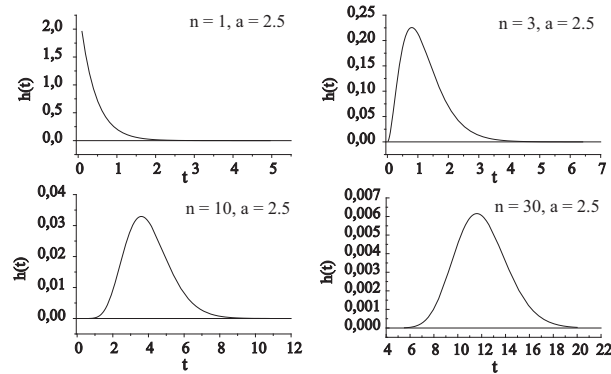
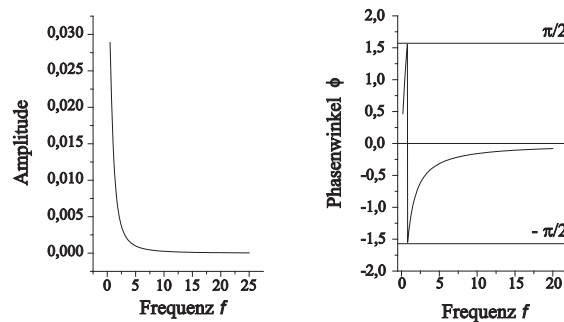


Abbildung A.7: Tiefpass 2-ter Ordnung, $b = -1.5$, $n = 1$.



Die folgenden Abbildungen A.7, A.8 und A.9. zeigen den Amplituden- und den Phasengang eines solchen Systems.

Über die verallgemeinerte Funktion δ läßt sich tatsächlich das Spektrum von $\sin$ – und $\cos$ – Funktionen anschreiben. Es ist

$$F(\sin \omega_0 x) = i\pi (\delta(\omega + \omega_0) - \delta(\omega - \omega_0)) \quad (\text{A.88})$$

$$F(\cos \omega_0 x) = \pi (\delta(\omega + \omega_0) + \delta(\omega - \omega_0)) \quad (\text{A.89})$$

Den Nachweis dieser Ausdrücke führt man, indem man für $\sin \omega_0 x$ und $\cos \omega_0 x$ die Beziehungen (A.19) und (A.20) einsetzt und (A.80) ausnutzt.

A.4.3 Eine Konstante

Es sei $f(x) = 1$, d.h. f sei konstant. Die Fouriertransformierte dieser Funktion ist dann durch

$$F(1) = \int_{-\infty}^{\infty} 1 \cdot e^{-i\omega x} dx = \int_{-\infty}^{\infty} e^{-i\omega x} dx$$

Abbildung A.8: Tiefpass 2-ter Ordnung, $b = -10.5$, $n = 2$.

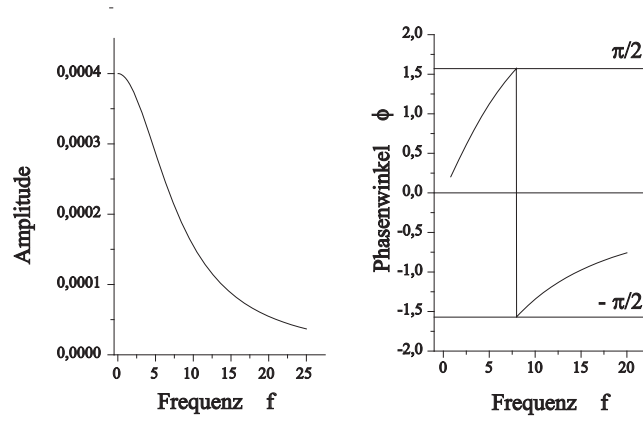
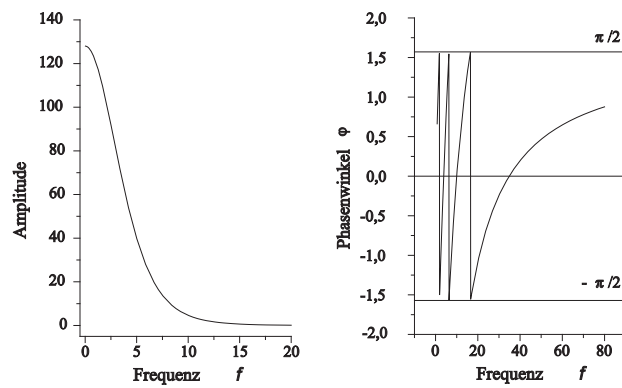


Abbildung A.9: Tiefpass 2-ter Ordnung, $b = -10.5$, $n = 7$.



gegeben. Andererseits ist nach (A.80)

$$\delta(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{i\omega x} dx$$

Dann ist

$$\begin{aligned} F(1) &= \int_{-\infty}^{\infty} e^{-i\omega x} dx \\ &= \int_{-\infty}^{\infty} e^{i(-\omega)x} dx \\ &= 2\pi\delta(-\omega) \end{aligned} \tag{A.90}$$

A.4.4 Die Funktion $1/x$

Es sei

$$f(x) = \frac{1}{x} \tag{A.91}$$

Sicherlich ist $f(-x) = -f(x)$, d.h. die Funktion ist ungerade. Nach (A.65) ist aber die Fouriertransformierte durch

$$F(\omega) = -i \int_{-\infty}^{\infty} \frac{\sin \omega x}{x} dx$$

gegeben, und von dem rechten Integral läßt sich zeigen, dass

$$\int_{-\infty}^{\infty} \frac{\sin \omega x}{x} dx = \begin{cases} -i\pi, & \omega > 0 \\ i\pi, & \omega < 0 \end{cases}$$

Also folgt

$$F(\omega) = \begin{cases} -i\pi, & \omega > 0 \\ i\pi, & \omega < 0 \end{cases} \tag{A.92}$$

Wendet man nun den Satz A.3.2 über die Symmetrie an, so erhält man weiter

$$F(i/\pi t) = -i \operatorname{sgn}(\omega) \tag{A.93}$$

A.4.5 Die Schrittfunktion

Es sei $U(x)$ die *Schrittfunktion*, d.h. es sei

$$U(x) = \begin{cases} 1, & x > 0 \\ 0, & x < 0 \end{cases} \tag{A.94}$$

Diese Funktion kann in der Form

$$U(x) = \frac{1}{2} + \frac{1}{2} \operatorname{sgn}(x), \quad \text{mit } \operatorname{sgn}(x) = \begin{cases} +1, & x > 0 \\ -1, & x < 0 \end{cases} \tag{A.95}$$

dargestellt werden. Dann ist

$$F(U(x)) = F\left(\frac{1}{2} + \frac{1}{2}\operatorname{sgn}(x)\right)$$

Nach Beispiel ?? hat aber die Konstante $1/2$ die Fouriertransformierte

$$F\left(\frac{1}{2}\right) = \frac{1}{2} \int_{-\infty}^{\infty} e^{-i\omega x} dx = \pi\delta(-\omega)$$

und nach Beispiel ?? erhält man weiter, dass

$$F(U(x)) = \pi\delta(\omega) + \frac{1}{i\omega} \quad (\text{A.96})$$

A.4.6 Der Balken

Es sei

$$f(x) = \begin{cases} 1, & -a < x < a \\ 0, & x \leq -a, \quad x \geq a \end{cases} \quad (\text{A.97})$$

Die Fouriertransformierte von f ist durch

$$F(\omega) = \frac{1}{2\pi} \int_{-\infty}^{\infty} f(x)e^{-i\omega x} dx = \frac{1}{2\pi} \int_{-a}^a f(x)e^{-i\omega x} dx$$

gegeben. Es ist

$$\int_{-a}^a f(x)e^{-i\omega x} dx = -\frac{1}{i\omega} e^{i\omega x} \Big|_{-a}^a = \frac{1}{i\omega} (e^{i\omega a} - e^{-i\omega a}),$$

so dass unter Berücksichtigung von (A.20)

$$F(\omega) = \frac{\sin \omega a}{\pi\omega} \quad (\text{A.98})$$

Nun sei umgekehrt

$$f(x) = \frac{\sin \omega_0 x}{\pi\omega_0}, \quad -\infty < x < \infty \quad (\text{A.99})$$

Die Fouriertransformierte von f ist jetzt nach dem Symmetrie-Theorem

$$F(\omega) = \begin{cases} 1, & -a < \omega < a \\ 0, & \omega \leq -a, \quad \omega \geq a \end{cases} \quad (\text{A.100})$$

Die Gauß-Funktion

Es sei nun

$$f(x) = \exp\left(-\frac{x^2}{2\sigma^2}\right) \quad (\text{A.101})$$

Es werde zur Vereinfachung $a = 1/2\sigma^2$ gesetzt. Die Fouriertransformierte ist dann durch

$$F(\omega) = \int_{-\infty}^{\infty} \exp(-a(x^2 + i\omega x/a)) dx$$

gegeben. Die Auswertung des Integrals liefert dann

$$F(\omega) = \sqrt{\frac{\pi}{a}} \exp\left(-\frac{\omega^2}{4a}\right). \quad (\text{A.102})$$

Die Fouriertransformierte einer Gauß-Funktion ist also wieder eine Gauß-Funktion, allerdings mit einer "Varianz", die proportional zu $1/\sigma^2$ ist.

A.4.7 Die Gabor-Funktion

Es werde nun die Funktion

$$g(x) = \exp\left(-\frac{x^2}{2\sigma^2}\right) \cos \omega_0 x \quad (\text{A.103})$$

betrachtet. Diese Funktion ist eine *Gabor-Funktion*; sie ist eine Kosinus-Funktion mit einer Gauß-Funktion als Enveloppe; hier wird allerdings nur der Spezialfall einer speziellen Rotation betrachtet, in Abschnitt ?? wird der allgemeine Fall diskutiert. Gesucht ist die Fouriertransformierte der Funktion (A.103).

Ein einfacher Weg, die Fouriertransformierte zu finden, besteht in der Anwendung des Modulations-Theorems (Satz ??). Nach (??) gilt ja

$$F(f(x) \cos \omega_0 x) = \frac{1}{2}(F(\omega + \omega_0) + F(\omega - \omega_0))$$

wobei F die Fouriertransformierte von f ist. Man substituiert hier für f die Funktion

$$f(x) = \exp\left(-\frac{x^2}{2\sigma^2}\right);$$

die Fouriertransformierte dieser Funktion ist in (??) gegeben. Man findet

$$F(\omega) = \frac{1}{2\sigma} \sqrt{\frac{\pi}{2}} \left(\exp\left(-\frac{\sigma^2(\omega + \omega_0)^2}{2}\right) + \exp\left(-\frac{\sigma^2(\omega - \omega_0)^2}{2}\right) \right) \quad (\text{A.104})$$

Analog findet man für die Gabor-Funktion $f(x) \sin \omega_0 x$, f wieder die Gaußfunktion, die Fouriertransformierte gemäß

$$F(f \sin(\omega_0 x)) = \frac{1}{2i}(F(\omega - \omega_0) - F(\omega + \omega_0)) = \frac{i}{2}(F(\omega + \omega_0) - F(\omega - \omega_0)) \quad (\text{A.105})$$

d.h.

$$F(\omega) = \frac{i}{2\sigma} \sqrt{\frac{\pi}{2}} \left(\exp\left(-\frac{\sigma^2(\omega + \omega_0)^2}{2}\right) - \exp\left(-\frac{\sigma^2(\omega - \omega_0)^2}{2}\right) \right) \quad (\text{A.106})$$

A.4.8 Hermite-Funktionen

Es werde die Funktion $f(x) = e^{-x^2}$ betrachtet; die n -te Ableitung von f ist dann $d^n e^{-x^2}/dx^n$. Es sei

$$H_n(x) = (-1)^n e^{x^2} \frac{d^n e^{-x^2}}{dx^n} \quad (\text{A.107})$$

H_n heißt *Hermite-Polynom*. Die ersten Hermite-Polynome sind

$$\begin{aligned} H_0(x) &= 1, & H_1(x) &= 2x, \\ H_2(x) &= 4x^2 - 2, & H_3(x) &= 8x^3 - 12x, \\ H_4(x) &= 16x^4 - 48x^2 + 12, \end{aligned} \quad (\text{A.108})$$

Die Funktionen

$$u_n(x) = H_n(x) \exp(-x^2/2) \quad (\text{A.109})$$

heißen *Hermite-Funktionen*. Die Hermite-Funktionen bilden ein *orthogonales Funktionensystem*; es gilt

$$\int_{-\infty}^{\infty} H_m(x) H_n(x) e^{-x^2} dx = \begin{cases} 0, & m \neq n \\ 2^n n! \sqrt{\pi}, & m = n \end{cases} \quad (\text{A.110})$$

und dementsprechend bilden die Funktionen

$$\phi_n(x) = \frac{H_n(x) e^{-x^2/2}}{\sqrt{2^n n! \sqrt{\pi}}} \quad (\text{A.111})$$

ein *normiertes Orthogonalsystem*³. Es ist also möglich, eine gegebene Funktion $f(x)$ in der Form

$$f(x) = \sum_{n=0}^{\infty} c_n \phi_n(x), \quad c_n = \langle f, \phi_n \rangle \quad (\text{A.112})$$

darzustellen.

Es sollen noch die Fouriertransformierten der Hermite-Funktionen hergeleitet werden. Dazu sei angemerkt, dass diese Funktionen der Differentialgleichung

$$u'' - x^2 u + \lambda u = 0, \quad \lambda \in \mathbb{R} \quad (\text{A.113})$$

genügen, wie man durch Einsetzen von ϕ_n für u sofort nachprüft. Bildet man in (A.113) die Fouriertransformierte, so erhält man die Dgl

$$F(u'') - F(x^2 u) + \lambda F(u) = 0 \quad (\text{A.114})$$

Aus (??) folgt dann

$$(i\omega)^2 F(u) - \frac{1}{-i^2} F((-ix)^2(u)) + \lambda F(u) = 0$$

³vergl. Courant, R. und Hilbert, D.: Methoden der Mathematischen Physik I, Springer-Verlag Berlin 1968

und also

$$-\omega^2 F(u) + F''(u) + \lambda F(u) = 0 \quad (\text{A.115})$$

Diese Gleichung hat aber die gleiche Form wie die für u . Dies bedeutet, dass der Ausdruck für die Fouriertransformierte von u die gleiche Form wie u selbst hat, d.h. proportional zu u ist:

$$F(\omega) = \sqrt{2\pi}(-i)^n H_n(\omega) e^{-\omega^2/2} \quad (\text{A.116})$$

Die Proportionalitätsfaktoren $\sqrt{2\pi}(-i)^n$ müssen noch gesondert hergeleitet werden.

Beispiel A.4.1 Es sei

$$f(x) = \frac{1}{x} \quad (\text{A.117})$$

Sicherlich ist $f(-x) = -f(x)$, d.h. die Funktion ist ungerade. Nach (??) ist aber die Fouriertransformierte durch

$$F(\omega) = -i \int_{-\infty}^{\infty} \frac{\sin \omega x}{x} dx$$

gegeben, und von dem rechten Integral läßt sich zeigen, dass

$$\int_{-\infty}^{\infty} \frac{\sin \omega x}{x} dx = \begin{cases} -i\pi, & \omega > 0 \\ i\pi, & \omega < 0 \end{cases}$$

Also folgt

$$F(\omega) = \begin{cases} -i\pi, & \omega > 0 \\ i\pi, & \omega < 0 \end{cases} \quad (\text{A.118})$$

Wendet man nun den Satz ?? über die Symmetrie an, so erhält man weiter

$$F(i/\pi t) = -i \operatorname{sgn}(\omega) \quad (\text{A.119})$$

□

A.5 Hilberttransformierte

Die folgende Definition erlaubt es, bestimmte Beziehungen zwischen dem Realteil $P(\omega)$ und dem Imaginärteil $Q(\omega)$ der Systemfunktion eines kausalen Systems herzuleiten.

Definition A.5.1 Es sei $f(t)$ eine reelle Funktion. Dann heißt

$$F_h(t) = \frac{1}{\pi} \int_{-\infty}^{\infty} \frac{f(\tau)}{\tau - t} d\tau \quad (\text{A.120})$$

die Hilberttransformierte von f . Die komplexe Funktion

$$\hat{f}(t) = f(t) - iF_h(t) \quad (\text{A.121})$$

heißt das zu $f(t)$ gehörende analytisches Signal; die Hilberttransformierte F_h heißt auch die Quadraturfunktion von $f(t)$.

Die Hilberttransformierte $F_h(t)$ ist offenbar die Faltung der Funktion $f(t)$ mit der Funktion $f_h(t) = -1/(\pi t)$. Aus Satz ??, p. ??, folgt dann

$$F(g) = G_1(\omega)G_2(\omega). \quad (\text{A.122})$$

Dementsprechend ergibt sich zwischen den Fouriertransformierten $F(i\omega)$ von $f(t)$ und der Fouriertransformierten $F(F_h)$ von F_h sofort die Beziehung

$$F(F_h) = F(i\omega)F(f_h). \quad (\text{A.123})$$

Die Fouriertransformierte $F(f_h)$ ist durch⁴

$$F(f_h) = i \operatorname{sgn}(\omega) \quad (\text{A.124})$$

mit

$$\operatorname{sgn}(\omega) = \begin{cases} +1, & \omega > 0 \\ -1, & \omega < 0 \end{cases}$$

gegeben. Bracewell (1983) weist darauf hin, dass die Faltung von $f(t)$ mit $f_h(t) = -1/\pi t$, d.h. die Hilberttransformation, einer "merkwürdigen" Art von Filterung entspricht: die Amplituden der Spektralkomponenten von $f(t)$ bleiben unverändert, aber die Phasen werden positiv oder negativ um den Betrag $\pi/2$ verändert, je nach dem Wert von $\operatorname{sgn} \omega$. Dies wird veranschaulicht, wenn man z.B. die Hilberttransformation von $\cos \omega t$ und $\sin \omega t$ betrachtet:

$$F_h(\cos \omega t) = -\sin \omega t, \quad F_h(\sin \omega t) = \cos \omega t. \quad (\text{A.125})$$

Es sei $S(x)$ die Heaviside-Funktion; $S(x) = 0$ für alle $x < 0$, und $S(x) = 1$ für $x \geq 0$. Es kann gezeigt werden, dass die Fouriertransformierte der Funktion $\frac{1}{2}\delta(t) + i/(2\pi t)$ gerade durch $S(\omega)$ gegeben ist, also

$$F\left(\frac{1}{2}\delta(t) + \frac{i}{2\pi t}\right) = \begin{cases} 1, & \omega \geq 0 \\ 0, & \omega < 0. \end{cases}$$

Es sei nun $f(t)$ eine reelle Funktion von t , und $F(\omega)$ sei die Fouriertransformierte von f . Nach Satz ?? ist

$$F(g_1 g_2) = G_1(\omega) * G_2(\omega), \quad (\text{A.126})$$

(*Frequenzfaltung*). Man betrachte nun das Produkt $2S(\omega)F(\omega)$. Sie ist das mit 2 multiplizierte Spektrum von f , bei dem aber alle negativen Frequenzen abgeschnitten werden. Gesucht ist die Funktion $\tilde{f}(t)$, deren Fouriertransformierte durch $2S(\omega)F(\omega)$ gegeben ist. Nach (A.122) muß $\tilde{f}$ durch 2 mal die Faltung der Funktion f und der inversen Transformierten von $S(\omega)$ sein, so dass

$$\tilde{f}(t) = 2 \left(\frac{1}{2}\delta(t) + \frac{i}{2\pi t} \right) * f(t),$$

⁴sgn: Abkürzung für signum = (Vor-)Zeichen.

d.h.

$$\tilde{f}(t) = 2 \cdot \frac{1}{2} \int_{-\infty}^{\infty} \delta(t - \tau) f(\tau) d\tau + \frac{i}{\pi} \int_{-\infty}^{\infty} \frac{f(\tau)}{t - \tau},$$

und demnach

$$\tilde{f}(t) = f(t) - iF_h(t). \quad (\text{A.127})$$

Dies bedeutet, dass $\tilde{f}(t) = \hat{f}(t)$, d.h. $\tilde{f}(t)$ ist gleich dem analytischen Signal. Die Spektren der Funktion $f(t)$ und des zugehörigen analytischen Signals stehen also in einer bestimmten Beziehung zueinander: ist $F(f) = F(\omega)$, so ist $F(\tilde{f}) = 2S(\omega)F(\omega)$.

A.6 Laplace-Transformierte

Die Bedingung (A.85) ist für eine Reihe von Funktionen, die für die Systemanalyse von Bedeutung sind, nicht erfüllt. So erfüllt schon $f(t) = e^{\alpha t}$, $\alpha > 0$, die Bedingung (A.85) nicht. Auch für die wichtige Schrittfunktion $f(t) = 0$, $t < 0$, $f(t) = 1$, $t \geq 0$ ist offenbar die Bedingung (A.85) nicht erfüllt; hier muß aber gesagt werden, dass es sich bei (A.85) nur um eine hinreichende, nicht um eine notwendige Bedingung handelt. Hsu (1967) leitet die Fouriertransformierte der Schrittfunktion her, und es zeigt sich, dass man nicht einfach (A.53) anwenden kann.

Für gegebene Funktion $f(t)$ betrachten wir nun die neue Funktion $g(t) := f(t)e^{-\alpha t}$. Für $g(t)$ ist die Bedingung der absoluten Integrierbarkeit (A.85) für die meisten in der Praxis vorkommenden Funktionen erfüllt. Weiter gelte $f(t) = 0$ für $t < 0$. Die Fouriertransformierte von $g(t)$ ist nun

$$H(s) = \int_0^{\infty} f(t)e^{-\alpha t - i\omega t} dt = \int_0^{\infty} f(t)e^{-st} dt, \quad s = \alpha + i\omega. \quad (\text{A.128})$$

Statt von der Fouriertransformierten von g spricht man nun von der (einseitigen) *Laplace-Transformierten* von f . Der Unterschied zur Fouriertransformierten von f ergibt sich dadurch, dass

- (i) das Integral von 0 bis ∞ bestimmt wird, nicht von $-\infty$ bis ∞ , und
- (ii) dass nicht das Integral von $f(t)e^{-i\omega t}$, sondern von $f(t)e^{st}$, s eine komplexe Zahl, gebildet wird.

Der Vorteil von (A.128) ist eben, dass $H(s)$ im allgemeinen auch dann noch existiert, wenn die Fouriertransformierte $H(i\omega)$ nicht existiert.

Beispiel A.6.1 (a) Gegeben sei $f(t) = e^{\alpha t}$, $\alpha > 0$. Gesucht ist die Laplace-Transformierte. Nach (A.128) ist $H(s) = \int e^{(\alpha-s)t} dt$, d.h.

$$L(e^{\alpha t}) = H(s) = \frac{1}{s - \alpha}, \quad \alpha. \quad (\text{A.129})$$

(b) Es sei $f(t) = t^k e^{\alpha t}$. Man findet

$$L(t^k e^{\alpha t}) = H(s) = \frac{k!}{(s - \alpha)^{k+1}}. \quad (\text{A.130})$$

(c) Für die Schrittfunktion findet man sofort

$$H(s) = \frac{1}{s}. \quad (\text{A.131})$$

□

Eine ausführliche Einführung in die Theorie der Laplace-Transformierten gibt Churchill (1958); dort findet man auch Tabellen mit den Transformationen der wichtigsten Funktionen, so dass man sich die Integration komplexer Funktionen weitgehend ersparen kann.

A.7 Die Schwartzsche Ungleichung

Nach der Schwartzschen Ungleichung gilt für irgend zwei Funktionen f_1 und f_2

$$\left| \int_a^b f_1(\xi) f_2(\xi) d\xi \right|^2 \leq \int_a^b |f_1(\xi)|^2 d\xi \int_a^b |f_2(\xi)|^2 d\xi, \quad (\text{A.132})$$

wobei f_1 und f_2 reell oder komplex sein können. Das Gleichheitszeichen gilt, wenn $f_2(\xi) = k f_1^*(\xi)$ für alle $\xi \in [a, b]$, k eine Konstante (ist f_1 reell, ist $f_1^* = f_1$). Die Aussage gilt gleichermaßen für Funktionen zweier Variablen, also $f_i(\xi, \eta)$, $i = 1, 2$.

Beweis: Es wird der allgemeine Fall komplexer Funktionen angenommen, da er den Spezialfall reeller Funktionen enthält. Nach (A.13), Seite 157, gilt für eine komplexe Zahl z die Darstellung $z = |z|e^{j\varphi}$. Im allgemeinen Fall sind die f_i komplex, also gilt für die komplexe Zahl

$$z = \int_a^b f_1(\xi) f_2(\xi) d\xi$$

die Darstellung

$$\int_a^b f_1(\xi) f_2(\xi) d\xi = \left| \int_a^b f_1(\xi) f_2(\xi) d\xi \right| e^{j\varphi}, \quad j = \sqrt{-1},$$

wobei φ der Phasenwinkel von z ist. Es sei nun $x \in \mathbb{R}$ eine beliebige Konstante. Dann gilt jedenfalls

$$\begin{aligned} \int_a^b |f_1 - x e^{j\varphi} f_2|^2 d\xi &= \int_a^b |f_1|^2 d\xi - x \left[e^{-j\varphi} \int_a^b f_1 f_2 d\xi + e^{j\varphi} \int_a^b f_1^* f_2^* d\xi \right] \\ &\quad + x^2 \int_a^b |f_2|^2 d\xi \\ &= \underbrace{\int_a^b |f_1|^2 d\xi}_{\alpha} + x^2 \underbrace{\int_a^b |f_2|^2 d\xi}_{\beta} - 2x \underbrace{\left| \int_a^b f_1 f_2 d\xi \right|}_{\gamma} \geq 0. \end{aligned}$$

Also

$$\int_a^b |f_1 - x e^{j\varphi} f_2|^2 d\xi = \alpha + x^2 \beta - 2x \gamma \geq 0,$$

oder, nach Division durch β

$$x^2 - 2x\frac{\gamma}{\beta} + \frac{\alpha}{\beta} \geq 0.$$

Nun ist

$$(x - \gamma/\beta)^2 = x^2 - 2\gamma/\beta + (\gamma/\beta)^2,$$

mithin läßt sich die Ungleichung in der Form

$$(x - \gamma/\beta)^2 - (\gamma/\beta)^2 + \alpha/\beta \geq 0$$

und dementsprechend in der Form

$$\frac{(x\beta - \gamma)^2}{\beta^2} - \frac{\gamma^2}{\beta^2} + \frac{\alpha\beta}{\beta^2} \geq 0$$

schreiben. Diese Ungleichung muß für beliebige x gelten; also hat man nach Multiplikation mit β^2 und $x = \gamma/\beta$

$$\alpha\beta \geq \gamma^2,$$

und dies war zu zeigen, □

A.8 2-dimensionale Signale

A.8.1 Allgemeine Definition

Ein Stimulismuster wird im Allgemeinen durch eine Funktion $f(x, y)$ zweier Ortskoordinaten x und y definiert; will man außerdem den zeitlichen Verlauf beschreiben, erhält man die Funktion $f(x, y, t)$, wenn t für die Zeit steht. Es soll hier die Fourier-Transformierte von $f(x, y)$ angegeben werden; die Verallgemeinerung auf drei Variablen ist dann direkt.

Die Fourier-Transformierte von f ist durch

$$F(u, v) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x, y) e^{-i(ux+vy)} dx dy \quad (\text{A.133})$$

gegeben. Dieser Ausdruck ergibt sich, wenn man nacheinander die Fouriertransformation bezüglich x und dann bezüglich y durchführt. Die inverse Transformation ist

$$f(x, y) = \frac{1}{4\pi^2} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} F(u, v) e^{iu+vy} du dv. \quad (\text{A.134})$$

Da die 2-dimensionale Fourier-Transformierte sich als Folge einfacher Fourier-Transformationen darstellen läßt, übertragen sich die Aussagen über diese Transformationen. Schreibt man

$$f(x, y) \Longleftrightarrow F(u, v) \quad (\text{A.135})$$

für ein Fourier-Paar von Funktionen, so lassen sich die Aussagen wie folgt zusammenfassen, wobei f^* und F^* für konjugiert komplexe Funktionen stehen:

$$f^*(x, y) \iff F^*(-u, -v) \quad (\text{A.136})$$

$$F(x, y) \iff 4\pi^2 f(-u, -v) \quad (\text{A.137})$$

$$f(x - x_0, y - y_0) \iff F(u, v)e^{-i(ux_0 + vy_0)} \quad (\text{A.138})$$

$$f(x, y)e^{i(ux_0 + vy_0)} \iff F(u - u_0, v - v_0) \quad (\text{A.139})$$

$$f(ax, by) \iff \frac{1}{|ab|} F\left(\frac{u}{a}, \frac{v}{b}\right) \quad (\text{A.140})$$

Eine Verallgemeinerung von (A.140) erweist sich als nützlich, wenn man nicht radial-symmetrische Stimuli in bestimmten Rotationen präsentieren will. Man findet

$$f(a_1x + b_1y, a_2x + b_2y) \iff \frac{1}{|a_1b_2 - a_2b_1|} F(A_1u + A_2v, B_1u + B_2v). \quad (\text{A.141})$$

Hierin sind $A_1, \dots, B_2$ die Elemente einer Matrix:

$$\begin{bmatrix} A_1 & B_1 \\ A_2 & B_2 \end{bmatrix} = \begin{bmatrix} a_1 & b_1 \\ a_2 & b_2 \end{bmatrix}^{-1}. \quad (\text{A.142})$$

Einen Spezialfall erhält man, wenn $f(x, y) = f_1(x)f_2(y)$. Dann gilt

$$f(x, y) = f_1(x)f_2(y) \iff F_1(u)F_2(v). \quad (\text{A.143})$$

Polarkoordinaten: Gelegentlich ist es nützlich, von Cartesischen zu Polarkoordinaten überzugehen. Dazu führt man die folgenden Beziehungen ein:

$$x = r \cos(\theta), \quad y = r \sin(\theta), \quad u = w \cos(\varphi), \quad v = w \sin(\varphi). \quad (\text{A.144})$$

Für das Fourier-Transformpaar f und F erhält man dann

$$f_0(ar, \theta + \theta_0) \iff \frac{1}{a^2} F_0\left(\frac{w}{a}, \varphi + \theta_0\right). \quad (\text{A.145})$$

Rotiert man also ein Muster um einen Winkel θ_0 , so wird die Fourier-Transformierte um den gleichen Winkel rotiert.

Man kann r als die Länge des Vektors von $(0, 0)$ zum Punkt (x, y) interpretieren, denn offenbar ist

$$\sqrt{x^2 + y^2} = \sqrt{r^2 \cos^2(\varphi) + r^2 \sin^2(\varphi)} = r \sqrt{\cos^2(\varphi) + \sin^2(\varphi)} = r,$$

denn $\cos^2(\varphi) + \sin^2(\varphi) = 1$. u und v tauchen in (A.133) als Frequenzkomponenten auf, und nach (A.144) findet man auf analoge Weise

$$\sqrt{u^2 + v^2} = \sqrt{w^2 \cos^2(\varphi) + w^2 \sin^2(\varphi)} = w. \quad (\text{A.146})$$

w ist die Frequenz in Bezug auf die Orientierung

$$\varphi = \tan^{-1}(v/u). \quad (\text{A.147})$$

yy

A.8.2 Gabor-Muster

Stochastische Prozesse und Verteilungen

A.9 Stochastische Prozesse

Es werden einige Grundbegriffe aus der Theorie der stochastischen Prozesse präsentiert.

Es werde zunächst an den Begriff der zufälligen Veränderlichen erinnert. Gegeben sei ein Wahrscheinlichkeitsraum (Ω, A, P) , wobei Ω ein Stichprobenraum, d.h. eine Menge von Elementarereignissen ist, A ist eine σ -Algebra und P ist ein Wahrscheinlichkeitsmaß. Dann heißt die Abbildung $X : \Omega \rightarrow \mathbb{R}$ eine (reelle) zufällige Veränderliche, wenn $P(X(\omega) \in B) = P(\omega \in X^{-1}(B))$ mit $X^{-1}(B) \in A$, $B \subseteq \mathbb{R}$. Die Verteilungsfunktion von X ist durch

$$F(x) = P(X \leq x) = P(X \in (-\infty, x]) \quad (\text{A.148})$$

gegeben. Die Dichte von X ist durch

$$f(x) = \frac{dF(x)}{dx} \quad (\text{A.149})$$

gegeben.

Definition A.9.1 Mit $X_t = \{X(t), t \in I\}$ werde eine Familie von zufälligen Veränderlichen bezeichnet, wobei I ein Intervall aus $\mathbb{R}$ ist. Die Familie X_t heißt stochastischer Prozeß, wenn die $X(t) \in X_t$ einen gemeinsamen Wahrscheinlichkeitsraum (Ω, A, P) haben. Für jedes $\omega \in \Omega$ heißt die Abbildung $t \rightarrow X(t, \omega)$ eine Realisierung oder Trajektorie von X_t . Ist die Menge $I \subseteq \mathbb{R}$ abzählbar, so heißt X_t stochastischer Prozeß mit diskreter Zeit, ist $I \subseteq \mathbb{R}$ überabzählbar, so heißt X_t stochastischer Prozeß mit stetiger Zeit.

Ist $X(t, \omega) \in \mathbb{R}^d$, $d \in \mathbb{N}$, so heißt X_t ein d -dimensionaler Prozeß. $X(t, \omega)$ repräsentiert dann einen Vektor zufälliger Funktionen.

Für festes $\omega \in \Omega$ ist $X(t, \omega)$ eine Funktion der Zeit. Diese Funktion ist zufällig insofern, als sie ebenfalls von dem eben zufällig auftretenden ω abhängt. Diese Abhängigkeit von ω bedeutet noch nicht, dass $X(t, \omega)$ irregulär aussieht, sie kann auch eine Gerade über einem bestimmten Intervall sein. Die Abhängigkeit von ω kann sich u. U. nur auf die Werte der Parameter der Geraden $X = at + b$, $t_1 \leq t \leq t_2$, auswirken, d.h. etwa $a = a(\omega)$, $b = b(\omega)$. Im allgemeinen wird man aber nicht nur die Menge der Geraden zulassen, sondern *alle* Funktionen über einem Intervall $[t_1, t_2]$, und damit eben auch die "zerknitterten". Man kann auch den Wert von t festhalten; dann ist $X(t, \omega)$ eine von ω abhängige zufällige Veränderliche. Welche Interpretation man für $X(t, \omega)$ wählt - als Funktion der Zeit für zufällig gewähltes ω oder X als von ω anhängige zufällige Veränderliche für gegebenes t - hängt im Prinzip nur davon

ab, welche der beiden Auffassungen die gerade vorliegende statistische Fragestellung leichter handhabbar macht.

Ist ein expliziter Bezug auf ω nicht notwendig, wird einfach $X(t)$ statt $X(t, \omega)$ geschrieben.

Es sei $I_n = \{t_1, \dots, t_n\}$ eine endliche Menge von Werten aus I und es sei $P(I_n)$ die gemeinsame Verteilung der zu den t_i , $i = 1, \dots, n$ korrespondierenden zufälligen Veränderlichen $X_1, \dots, X_n$. Weiter sei $\{I_n\}$ die Menge der möglichen Mengen von n (Zeit-)Punkten aus I . Dann heißt $\{P(I_n)\}$ die *Familie der endlich-dimensionalen Verteilungen* des stochastischen Prozesses X_t . Viele den stochastischen Prozess beschreibenden statistischen Größen lassen sich über diese Familie bestimmen. Die wichtigsten dieser statistischen Größen werden in der folgenden Definition eingeführt:

Definition A.9.2 Für festes $t \in I$ ist $F(x; t) = P(X(t) \leq x)$ die Verteilungsfunktion von $X(t)$, und $f(x; t) = dF(x; t)/dx$ die dazu korrespondierende Dichtefunktion. Dann ist

$$\mu(t) = E(X(t)) = \int_{-\infty}^{\infty} x f(x; t) dx \quad (\text{A.150})$$

die Mittelwertsfunktion von X_t , und

$$R(t_1, t_2) = E[X(t_1), X(t_2)] = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_1 x_2 f(x_1, x_2; t_1, t_2) dx_1 dx_2 \quad (\text{A.151})$$

die Autokorrelationsfunktion von X_t ; dabei ist $f(x_1, x_2; t_1, t_2)$ die gemeinsame Dichte von X_1, X_2 . Die Größe

$$Kov(t_1, t_2) = E[(X(t_1) - \mu(t_1))(X(t_2) - \mu(t_2))] = R(t_1, t_2) - \mu(t_1)\mu(t_2) \quad (\text{A.152})$$

heißt Kovarianzfunktion. Dementsprechend ist die Varianz von X_t durch

$$\sigma^2(t) = Kov(t, t) = R(t, t) - \mu^2(t) \quad (\text{A.153})$$

erklärt.

Für $t_1 = t_2$ nimmt $R(t_1, t_2)$ ein Maximum an. Um dies zu sehen, betrachte man

$$E[(X(t_1) - X(t_2))^2] = E[X^2(t_1)] + E[X^2(t_2)] - 2E[X(t_1)X(t_2)] \geq 0$$

Nun ist aber $E[X^2(t_1)] = R(t_1, t_1)$, $E[X^2(t_2)] = R(t_2, t_2)$, $E[X(t_1), X(t_2)] = R(t_1, t_2)$, und es folgt

$$R(t_1, t_1) + R(t_2, t_2) \geq 2R(t_1, t_2) \quad (\text{A.154})$$

für alle $t_1 \neq t_2$. Für $t_2 \rightarrow t_1$ bzw. $t_2 \rightarrow t_1$ folgt dann die Gleichheit der beiden Ausdrücke, und dies heißt ja, dass $R(t_1, t_2)$ nicht größer als $R(t_1, t_1)$ oder $R(t_2, t_2)$ werden kann.

Definition A.9.3 Es seien X_t und Y_t zwei stochastische Prozesse; dann heißt

$$R_{xy}(t_1, t_2) = E[X(t_1)Y(t_2)] \quad (\text{A.155})$$

Die Kreuzkorrelation der Prozesse X_t und Y_t .

Definition A.9.4 Es sei $X_t = \{X(t), t \in (-\infty, \infty)\}$ ein stochastischer Prozess. X_t heißt stationär, wenn für jede endliche Menge $I_n = \{t_1, \dots, t_n\}$ die gemeinsame Verteilung von $\{X_{t_1+t_0}, \dots, X_{t_n+t_0}\}$ von t_0 unabhängig ist. X_t heißt stationär im weiten Sinne, wenn gilt: (i) $E(X^2(t)) < \infty$, (ii) $\mu = E(X(t)) = \text{konstant}$, (iii) $Kov(s, t)$ hängt nur von der Differenz $t - s$ ab.

Für die Autokorrelation $R(t_1, t_2)$ bedeutet Stationarität und Stationarität im weiten Sinne, dass $R(t_1, t_2) = R(t_2 - t_1) = R(\tau)$ gilt, d.h. R hängt nur von der Differenz $\tau = t_2 - t_1$, nicht aber von den Werten t_1 und t_2 ab. Die Beziehung (A.154) läßt sich dann zu

$$R(0) = \max_{\tau} R(\tau) \quad (\text{A.156})$$

spezifizieren.

Definition A.9.5 Der stochastische Prozess $X_t = \{X(t), t \in I\}$ sei stationär. Lassen sich alle Statistiken des Prozesses (Mittelwertfunktion, Autokorrelationsfunktion, etc.) bereits aus einer Trajektorie $X(t, \omega)$ berechnen, so heißt der Prozeß ergodisch.

Die Ergodizität eines Prozesses ist für empirische Untersuchungen von großer Bedeutung, da man eben im allgemeinen *nur eine* Trajektorie des Prozesses beobachten und deshalb nur sie zur Berechnung von Statistiken heranziehen kann.

Die Eigenschaften eines stochastischen Prozesses lassen sich weiter durch Bezug auf den Energiebegriff beschreiben. Um diesen Zusammenhang zu verdeutlichen, sei daran erinnert, dass viele Funktionen sich als Überlagerung von Sinus-Funktionen dargestellt werden können; dazu muß nur die Amplitude $|F(\omega)|$ und die Phase $\varphi(\omega)$ für jede Frequenz ω geeignet gewählt werden; wie im Anhang hergeleitet wird, können diese beiden Größen durch Berechnung der Fourier-Transformierten $F(\omega)$ bestimmt werden. Wir nehmen an, dass die Fourier-Transformierte auch für jede Trajektorie $X(t)$ des stochastischen Prozesses X_t berechnet werden kann, dass also

$$F(\omega) = \int_{-\infty}^{\infty} X(t) e^{-i\omega t} dt, \quad (\text{A.157})$$

$i = \sqrt{-1}$, existiert. $F(\omega)$ ist im allgemeinen eine komplexe Zahl, $F(\omega) = F_1(\omega) + iF_2(\omega)$, wobei F_1 der Real- und F_2 der Imaginärteil von F ist. F_1 und F_2 sind jeweils reelle Funktionen von ω . Die Amplitude für die Frequenz ω ergibt sich gemäß $|F(\omega)| = \sqrt{F_1^2(\omega) + F_2^2(\omega)}$, und $\varphi(\omega) = \tan^{-1}[F_2(\omega)/F_1(\omega)]$.

Weiter kann man jeder Trajektorie $X(t)$ entsprechend bestimmter Begriffsbildungen in der Physik eine "Energie" zuordnen, nämlich

$$E = \int_{-\infty}^{\infty} |X(t)|^2 dt, \quad (\text{A.158})$$

Ist $F(\omega)$ die zu $X(t)$ gehörige Fourier-Transformierte, so gilt *Parsevals Theorem*, demzufolge

$$E = \int_{-\infty}^{\infty} |X(t)|^2 dt = \int_{-\infty}^{\infty} |F(\omega)|^2 d\omega. \quad (\text{A.159})$$

Da $X(t) = X(t, \omega)$ eine zufällig gewählte Funktion ist, sind die Größen $F(\omega)$, $|F(\omega)|$, $\varphi(\omega)$ und E zufällige Veränderliche und haben deshalb möglicherweise einen Erwartungswert. Man ist insbesondere an der durchschnittlichen Energie der Trajektorien und damit des Prozesses X_t interessiert. Es zeigt sich nun, dass sich die durchschnittliche Energie des Prozesses als Fourier-Transformierte der Autokorrelation $R(\tau)$ (bei einem stationären Prozeß) darstellen läßt. Dies sieht man wie folgt:

Es sei $\bar{F}$ die zu F konjugiert komplexe Zahl, d.h. es sei $\bar{F} = F_1(\omega) - iF_2(\omega)$; dann ist $|F(\omega)|^2 = F\bar{F}$. Ebenso kann man $X(t)$ vorübergehend als komplexe Funktion auffassen und $X\bar{X}$ für $|X(t)|^2$ schreiben. Nun ist $|F(\omega)|^2$ durch

$$|F(\omega)|^2 = F(\omega)\bar{F}(\omega) = \left| \int_{-\infty}^{\infty} X(t)e^{-i\omega t} dt \right|^2 = \int_{-\infty}^{\infty} X(u)e^{-i\omega u} du \int_{-\infty}^{\infty} \bar{X}(v)e^{i\omega v} dv$$

gegeben (dabei ist die Integrationsvariable t formal durch u und v ersetzt worden, um Verwechslungen zu vermeiden). Dann gilt aber

$$|F(\omega)|^2 = \int_{-\infty}^{\infty} X(u)e^{-i\omega u} du \int_{-\infty}^{\infty} X^*(v)e^{i\omega v} dv = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} X(u)X^*(v)e^{-i\omega(u-v)} dudv.$$

Da $F(\omega)$ eine zufällige Veränderliche ist, muß auch $|F(\omega)|^2$ eine zufällige Veränderliche sein, mithin kann man den Erwartungswert von $|F(\omega)|^2$ bilden, und da $e^{-i\omega(u-v)}$ keine zufällige Veränderliche ist, folgt

$$E(|F(\omega)|^2) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} E(X(u)\bar{X}(v))e^{-i\omega(u-v)} dudv. \quad (\text{A.160})$$

Aber $E(X(u)\bar{X}(v)) = R(u-v)$, R die Autokorrelation eines stationären Prozesses (vergl. Definition A.9.4). Diese Beziehung sollte gezeigt werden. Damit hat man

Definition A.9.6 *Es sei X_t ein stationärer Prozeß und $F(\omega)$ seien die Fourier-Transformierten der Trajektorien $X(t)$ von X_t . Dann heißt*

$$S(\omega) = E(|F(\omega)|^2) = \int_{-\infty}^{\infty} R(t)e^{-i\omega t} dt \quad (\text{A.161})$$

die Spektraldichtefunktion des stochastischen Prozesses X_t .

Die Spektraldichtefunktion ist also die Fourier-Transformierte der Autokorrelationsfunktion. Dann kann man $R(\tau)$ als inverse Transformation von $S(\omega)$ darstellen:

$$R(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} S(\omega)e^{i\omega t} d\omega. \quad (\text{A.162})$$

Aus der Definition von R folgt aber, dass $R(0) = \sigma^2$ die Varianz des Prozesses X_t , und gemäß (A.162) ist dann

$$\sigma^2 = R(0) = \frac{1}{2\pi} \int_{-\infty}^{\infty} S(\omega) d\omega; \quad (\text{A.163})$$

σ^2 gleicht also der *durchschnittlichen Energie* des Prozesses (vergl (A.159) und (A.161), Seite 188).

Es sei X_t ein stationärer Prozeß mit der Autokorrelationsfunktion

$$R(\tau) = \begin{cases} \sigma^2(\tau) = \sigma^2, & \tau = 0 \\ 0, & \tau \neq 0 \end{cases} \quad (\text{A.164})$$

Dies bedeutet, dass beliebig benachbarte Werte von $X(t)$ stochastisch unabhängig und damit unkorreliert sind. Man überlegt sich leicht, dass ein solcher Prozeß in der Natur kaum auftreten kann, denn jeder an die Gesetze der Physik gekoppelte Prozeß benötigt Zeit für eine Veränderung, und die Forderung, dass die Werte $X(t)$ und $X(t + \epsilon)$ für $\epsilon > 0$ *beliebig* klein unkorreliert sind für alle $t \in I$ bedeutet ja, dass die Trajektorien von X_t sich beliebig schnell verändern können. Soweit psychische Prozesse eine physikalische Basis haben gelten diese Bemerkungen auch für psychische Prozesse. Prozesse mit der Autokorrelation (A.164) heißen dann auch *generalisierte Prozesse*. Es zeigt sich aber, dass generalisierte Prozesse häufig gut als erste Annäherung an die tatsächlichen Prozesse gewählt werden können; ihre angenehmen mathematischen Eigenschaften bedeuten dabei eine wesentliche Vereinfachung des formalen Aufwandes bei der Diskussion der Daten. Dazu werde die zu (A.164) gehörende Spektraldichtefunktion betrachtet. Es ist

$$S(\omega) = \sigma^2 \int_{-\infty}^{\infty} R(t) e^{-i\omega t} dt = \sigma^2, \quad (\text{A.165})$$

d.h. aber $S(\omega) = \sigma^2$ ist eine konstante für alle ω . Dies bedeutet, dass jede Frequenz ω gleich viel Energie zur durchschnittlichen Energie beiträgt. Hat man eine Lichtquelle, in der jede Wellenlänge (Frequenz) mit gleicher Energie vorkommt, so erhält man bekanntlich weißes Licht. In Analogie hierzu spricht man dann auch von X_t als von einem *weißen Rauschen*; das Wort *Rauschen* erinnert dabei an das durch Zufallsprozesse im Radio oder Telefonhörer auftretende Rauschen, und es ist *weiß*, weil es wie das weiße Licht alle Frequenzen mit gleicher Energie enthält. Die Annahme weißen Rauschens ist überall dort angebracht, wo die Spektraldichte über einem hinreichend breiten Frequenzband angenähert konstant ist.

Der Prozeß des Weißen Rauschens ist ein Spezialfall, der nicht für alle Prozesse als angenäherte Beschreibung gewählt werden kann. Im allgemeinen sind benachbarte Werte $X(t)$ und $X(t + \Delta t)$ in einem von Δt abhängendem Ausmaß korreliert, d.h. die Werte $X(t)$ und $X(t + \Delta t)$ sind stochastisch abhängig. Man kann nun, zur weiteren Diskussion der Abhängigkeiten zwischen Werten von $X(t)$ zu verschiedenen Zeitpunkten $t_1, \dots, t_n$, vom Begriff der bedingten Verteilung Gebrauch machen. Dazu faßt man $X_1 = X(t_1), X_2 = X(t_2), \dots, X_n = X(t_n)$ als zufällige Veränderliche auf. Die Abhängigkeiten lassen sich dann über die bedingten Dichten

$$f(x_1, \dots, x_k | x_{k+1}, \dots, x_n) = \frac{f(x_1, \dots, x_n)}{f(x_{k+1}, \dots, x_n)} \quad (\text{A.166})$$

behandeln.

Beim Weißen Rauschen sind die $x_n, \dots, x_1$ *alle* stochastisch unabhängig voneinander, so lange nur die $t_1, \dots, t_n$ verschieden sind. Das Weiße Rauschen fluktuiert

also mit der höchstmöglichen Geschwindigkeit. Bei Prozessen mit Abhängigkeiten zwischen den x_i ist die Geschwindigkeit der Fluktuationen dann geringer. Es liegt nahe, ein Maß für diese Geschwindigkeit zu suchen; intuitiv ist klar, dass ein solches Maß mit der Autokorrelation zusammenhängen muß, da diese ja die Korrelationen der $X(t)$ für benachbarte Zeitpunkte angibt. Ein solches Maß kann wie folgt konstruiert werden.

Es seien $X(t)$ die Trajektorien des Prozesses X_t , und es werde angenommen, dass sie zumindest fast überall differenzierbar sind. Für eine spezielle Trajektorie sei dann

$$Y(t) = \frac{dX(t)}{dt} = X'(t); \quad (\text{A.167})$$

$Y(t)$ ist wieder eine Funktion von t , und da $X(t)$ eine zufällige Funktion ist, muß auch $Y(t)$ eine zufällige Funktion sein. Ist also $X_t = \{X(t), t \in I\}$ ein stochastischer Prozeß mit differenzierbaren Trajektorien, so ist $Y_t = \{Y(t), t \in I\}$ ebenfalls ein stochastischer Prozeß, dessen Trajektorien eben die Differentiale der Trajektorien $X(t)$ sind.

Wir suchen nun die Autokorrelationsfunktion des Prozesses Y_t . die Autokorrelation von X_t ist $R(t_1, t_2) = E[X(t_1)X(t_2)]$. Die Kreuzkorrelation der Prozesse X_t und Y_t ist dann $R_{xy} = E[X(t_1)X'(t_2)]$. Nun ist

$$E \left[X(t_1) \frac{X(t_2 + \epsilon) - X(t_2)}{\epsilon} \right] = \frac{R_{xx}(t_1, t_2 + \epsilon) - R_{xx}(t_1, t_2)}{\epsilon},$$

und für $\epsilon \rightarrow 0$ wird hieraus

$$R_{xy}(t_1, t_2) = \frac{\partial R_{xx}(t_1, t_2)}{\partial t_2}. \quad (\text{A.168})$$

Ebenso kann man

$$E \left[\frac{X(t_1 + \epsilon) - X(t_1)}{\epsilon} X'(t_2) \right] = \frac{R_{xy}(t_1 + \epsilon, t_2) - R_{xy}(t_1, t_2)}{\epsilon}$$

bilden, und für $\epsilon \rightarrow 0$ folgt hieraus

$$R_{yx}(t_1, t_2) = \frac{\partial R_{xy}(t_1, t_2)}{\partial t_1}. \quad (\text{A.169})$$

Aus (A.168) und (A.169) folgt dann

$$R_{yy}(t_1, t_2) = \frac{\partial^2 R_{xx}(t_1, t_2)}{\partial t_1 \partial t_2}; \quad (\text{A.170})$$

damit ist die Autokorrelation des Prozesses $Y_t = \{X'(t)|t \in I\}$ aus der Autokorrelation des Prozesses $X_t = \{X(t)|t \in I\}$ durch zweifache partielle Differentiation ableitbar. Insbesondere sei nun der Prozeß X_t stationär, dann gilt ja $R(t_1, t_2) = R(\tau)$, $\tau = t_2 - t_1$. Dann folgt insbesondere

$$R_{yy}(\tau) = -\frac{d^2 R(\tau)}{d^2 \tau}. \quad (\text{A.171})$$

Für $\tau = 0$ folgt dann

$$R_{yy}(0) = -\frac{d^2 R(0)}{d^2 \tau} = E[Y^2(t)] = E[(X'(t))^2], \quad (\text{A.172})$$

d.h. $R_{yy}(0)$ ist die Varianz σ_y^2 des Prozesses $Y_t = \{X'(t) | t \in I\}$. Ist σ_y^2 groß, so heißt dies, dass die *Varianz* der Veränderungen $X'(t)$ der Trajektorien des Prozesses X_t groß ist, dass die Trajektorien $X(t)$ des Prozesses X_t also stark fluktuieren bzw. dass die $X(t)$ keinen sehr regelmäßigen Verlauf haben. Die Spektraldichte von Y_t ist nun wieder durch die Fourier-Transformierte von R_{yy} gegeben, und die *durchschnittliche Energie des Prozesses* Y_t ist durch $R_{yy}(0)$ gegeben. Gemäß (A.172) ist aber $R_{yy}(0)$ auch mit Eigenschaften des Prozesses X_t verbunden. Diese Verbindung soll noch etwas elaboriert werden.

Es sei X_t stationär; dann ist $R(t_1, t_2) = R(\tau)$ mit $\tau = t_2 - t_1$. R ist maximal für $\tau = 0$, und weiter gilt

$$R(-\tau) = R(\tau); \quad (\text{A.173})$$

dies folgt leicht aus der Stationarität. Nun ist einerseits

$$S(\omega) = \int R(\tau) e^{-i\omega\tau} d\omega / 2\pi,$$

so dass umgekehrt

$$R(\tau) = \int S(\omega) e^{i\omega\tau} d\omega / 2\pi$$

geschrieben werden kann. Dann ist andererseits

$$R'(\tau) = \frac{dR(\tau)}{d\tau} = \int S(\omega) e^{i\omega\tau} d\omega / 2\pi$$

und

$$R''(\tau) = \frac{d^2 R(\tau)}{d\tau^2} = \int (i\omega)^2 S(\omega) e^{i\omega\tau} d\omega / 2\pi = - \int \omega^2 S(\omega) e^{i\omega\tau} d\omega / 2\pi,$$

so dass

$$\left. \frac{d^2(R(\tau))}{d\tau^2} \right|_{\tau=0} = -\frac{1}{2\pi} \int_{-\infty}^{\infty} \omega^2 S(\omega) d\omega \quad (\text{A.174})$$

folgt. Das Integral kann nun als zweites Moment der Spektraldichte $S(\omega)$ des Prozesses X_t aufgefaßt werden. Dies führt zu der

Definition A.9.7 *Es sei X_t ein stochastischer Prozess mit differenzierbaren Trajektorien, und es sei $Y_t = \{X'(t) | X'(t) = dX(t)/dt, t \in I\}$. Dann heißt*

$$\lambda_2 := E[(X'(t))^2] = - \left. \frac{d^2 R_{xx}(\tau)}{d\tau^2} \right|_{\tau=0} = -R''(0) \quad (\text{A.175})$$

das zweite Spektralmoment des stochastischen Prozesses X_t .

Natürlich ist $\lambda_2 = \sigma_y^2$. Notwendige Bedingung dafür, dass das zweite Spektralmoment λ_2 existiert, ist die Differenzierbarkeit von $R(\tau)$ an der Stelle $\tau = 0$. In Zusammenhang mit dynamischen Modellen werden allerdings oft spezielle Prozesse, die Diffusionsprozesse diskutiert, die *nirgends differenzierbare* Trajektorien haben. Man betrachtet hier einen anderen Begriff als das zweite Spektralmoment:

Definition A.9.8 *Es sei X_t ein stationärer Prozeß mit nicht notwendig differenzierbaren Trajektorien und der Autokorrelationsfunktion $R(t)$. Dann heißt*

$$\tau_{cor} = \frac{1}{R(0)} \int_0^\infty R(\tau) d\tau \quad (\text{A.176})$$

die Korrelationszeit des Prozesses.

Beispiel A.9.1 Gegeben sei ein stationärer stochastischer Prozeß X_t mit der Autokorrelationsfunktion $R(\tau) = \exp(-\alpha\tau^2)$, $\alpha > 0$. Für große α geht $R(\tau)$ schnell gegen Null, d.h. die Trajektorien fluktuieren schnell; für kleines α fluktuieren sie langsam. Es ist $dR(\tau)/d\tau = -2\alpha\tau R(\tau)$, und $d^2R(\tau)/d\tau^2 = -2\alpha R(\tau) + 4\alpha^2\tau^2 R(\tau)$. Dann ist $d^2R(\tau)/d\tau^2|_{\tau=0} = -2\alpha R(0)$. Aber $R(0) = 1$, so dass das zweite Spektralmoment durch $\lambda = 2\alpha$ gegeben ist.

Nun sei X_t ein Prozeß mit nicht differenzierbaren Trajektorien und der Autokorrelationsfunktion $R(\tau) = \sigma^2 e^{-\alpha\tau}$, $\alpha > 0$. Die Korrelationszeit ist

$$\tau_{cor} = \frac{1}{R(0)} \int_0^\infty R(\tau) d\tau = \frac{1}{\sigma^2} \int_0^\infty e^{-\alpha\tau} d\tau = \frac{1}{\sigma^2} \frac{1}{\alpha}.$$

Wieder gilt, dass τ_{cor} um so kleiner ist, je größer α ist. Will man das zweite Spektralmoment λ berechnen, so muß die zweite Ableitung bestimmt werden. Es ist $dR(\tau)/d\tau = -\alpha R(\tau)$, $d^2R(\tau)/d\tau^2 = \alpha^2 R(\tau)$ und $d^2R(\tau)/d\tau^2|_{\tau=0} = \sigma^2 \alpha^2 > 0$. Es ist aber

$$\lambda = - \left. \frac{d^2R(\tau)}{d\tau^2} \right|_{\tau=0},$$

so dass $\lambda < 0$ wäre. Offenbar macht die Berechnung des zweiten Spektralmoments hier keinen Sinn. $\square$

In (A.155) war die bedingte Dichte

$$f(x_1, \dots, x_k | x_{k+1}, \dots, x_n) = \frac{f(x_1, \dots, x_n)}{f(x_{k+1}, \dots, x_n)}$$

eingeführt worden. Hieraus ergibt sich ein für die Anwendung der Theorie stochastischer Prozesse wichtiger Spezialfall, der durch eine eigene Definition charakterisiert werden soll:

Definition A.9.9 *Es sei X_t ein stochastischer Prozeß und es gelte für jedes n und $t_1 < t_2 < \dots < t_n$*

$$P(x(t_n) \leq x_n | x(t_{n-1}), \dots, x(t_1)) = P(x(t_n) \leq x_n | x(t_{n-1})). \quad (\text{A.177})$$

Dann heißt X_t Markov-Prozeß.

Insbesondere Markov-Prozesse lassen sich oft durch Übergangswahrscheinlichkeiten charakterisieren. Es sei $p(x, t; x_0, t_0)$ die Wahrscheinlichkeit, dass der Prozeß zur Zeit t den Wert (Zustand) $x = x(t)$ annimmt, wenn er zur Zeit $t_0 \leq t$ den Wert (Zustand) $x_0 = x(t_0)$ angenommen hat. $p(x, t; x_0, t_0)$ heißt auch Übergangswahrscheinlichkeit. $p(x, t; x_0, t_0)$ ist eine bedingte Dichtefunktion. Es gilt sicherlich

$$p(x, t; x_0, t_0) \rightarrow \delta(x - x_0) \text{ für } t \rightarrow t_0, \quad (\text{A.178})$$

wobei δ die Dirac-Delta-Funktion ist, d.h. es ist $\delta(x - x_0) = 0$ für $x \neq x_0$ und $\delta(x - x_0) = \infty$ für $x = x_0$ unter der Nebenbedingung $\int \delta(x - x_0) dx = 1$. Da $p(x, t; x_0, t_0)$ eine Dichtefunktion ist, muß darüber hinaus

$$\int_{-\infty}^{\infty} p(x, t; x_0, t_0) dx = 1 \quad (\text{A.179})$$

gelten.

Es sei $t_0 < t_1 < t$. Dann gilt die *Chapman-Kolmogorov-Gleichung*:

$$p(x, t; x_0, t_0) = \int_{-\infty}^{\infty} p(x, t; x_1, t_1) p(x_1, t_1; x_0, t_0) dx_1. \quad (\text{A.180})$$

Die Gleichung ist eine direkte Konsequenz von (A.176). Denn sei $t_0 < t_1 < t_2$, und sei $p(x(t_2), x(t_1), x(t_0))$ die Wahrscheinlichkeit, dass der Prozeß zu den Zeiten t_0 , t_1 und t_2 die Werte $x(t_0)$, $x(t_1)$ und $x(t_2)$ annimmt. Nach dem Satz über bedingte Wahrscheinlichkeiten gilt dann

$$p(x(t_2), x(t_1), x(t_0)) = p(x(t_2)|x(t_1), x(t_0)) p(x(t_1)|x(t_0)).$$

Aber die Markov-Eigenschaft bedeutet, dass $p(x(t_2)|x(t_1), x(t_0)) = p(x(t_2)|x(t_1))$ und mithin

$$p(x_2, t_2; x_1, t_1; x_0, t_0) = p(x_2, t_2; x_1, t_1) p(x_1, t_1; x_0, t_0).$$

Andererseits ist

$$p(x_2, t_2; x_0, t_0) = \int_{-\infty}^{\infty} p(x_2, t_2; x_1, t_1; x_0, t_0) dx_1;$$

setzt man den vorangehenden Ausdruck für den Integranden ein, so erhält man die Chapman-Kolmogoroff-Gleichung.

Die Markov-Annahme unterliegt auch der Beschreibung dynamischer Systeme durch deterministische Differentialgleichungen. Ist nämlich $x(t) = (x_1(t), \dots, x_n(t))^T$ der Zustandsvektor eines solchen Systems, so ist es im allgemeinen durch eine Differentialgleichung der Form $\dot{x} = f(x, t)$ beschreibbar, wobei f ein geeignet gewählter Vektor von Funktionen ist. Die Beschreibung durch eine solche Differentialgleichung besagt, dass die Veränderung $\dot{x}$ des Zustandsvektors durch den *momentanen* Wert $x(t)$ und den einer eventuell auf das System einwirkenden "Stör"funktion ist. Der Markov-Prozeß kann als Verallgemeinerung einer solchen Konzeption von Systemen für den Fall, dass zufällige Effekte berücksichtigt werden müssen, aufgefaßt werden.

Die Annahme allerdings, dass Veränderungen momentan wirken und vergangene Zeitpunkte keine Rolle spielen sollen, kann genau genommen nur eine Approximation sein, da jede Veränderung Zeit benötigt. Die Markov-Annahme stellt deshalb eine Idealisierung dar, deren Vorteil eine erhebliche Erleichterung der mathematischen Analyse der jeweils interessierenden Prozesse ist. Der Nachteil, eben nur eine Approximation an die Wirklichkeit zu sein, spielt praktisch keine Rolle, wenn der Einfluß der Vergangenheit im Vergleich zu der interessierenden Zeitskala verschwindend klein ist.

A.10 Die Weibull-Verteilung

Definition A.10.1 Gegeben sei die zufälligen Veränderlichen X mit $-\infty < X \leq \eta_0$ und Y mit $\eta_0 \leq Y < \infty$. Die Verteilungsfunktionen von X bzw. Y seien durch

$$F_X(x) = P(X \leq x) = \begin{cases} \exp \left[- \left(\frac{\eta_0 - x}{a} \right)^\beta \right], & x \leq \eta_0. \\ 1, & x > \eta_0 \end{cases} \quad (\text{A.181})$$

$$F_Y(y) = P(Y \leq y) = \begin{cases} 1 - \exp \left[- \left(\frac{y - \eta_0}{a} \right)^\beta \right], & y \geq \eta_0 \\ 0, & y < \eta_0 \end{cases} \quad (\text{A.182})$$

gegeben. Dann heißen X bzw. Y Weibull-verteilt mit den Parametern η_0 , $a > 0$, und $\beta > 0$.

Anmerkungen:

1. Der Name Weibull-Verteilung geht auf den schwedischen Ingenieur Ernst Hjalmar Waloddi Weibull (1887 – 1979) zurück, der diese Verteilung im Zusammenhang mit Fragen der Bruchfestigkeit spröden Materials diskutierte. Fellers (1966), p. 54 Anmerkung, dass insbesondere die Verteilung (A.182) in Reliabilitätsuntersuchungen "rather mysteriously" unter dem Namen Weibull-Verteilung auftaucht, ist vermutlich damit zu erklären, dass die Verteilung schon vor Weibulls Arbeiten bekannt war.
2. Die zufälligen Veränderlichen X und Y sind durch die Beziehung $Y = 2\eta_0 - X$ miteinander verbunden. Ist die Verteilungsfunktion von X durch (A.181) gegeben, so folgt die Verteilungsfunktion (A.182) für Y , denn

$$P(Y \leq y) = 1 - P(X \leq 2\eta_0 - y) = 1 - \exp \left[- \left(\frac{y - \eta_0}{a} \right)^\beta \right].$$

3. Der Hintergrund für die Definition der Weibull-Verteilung durch (A.181) und (A.182) ist, dass beide Formen als Definition Weibull-Verteilung in der Literatur vorkommen. Die Form (A.181) wird in Johnson, Kotz und Balakrishna

(1994) und Kotz und Nadarajah (2000) als Verteilung vom Weibull-Typ bezeichnet. Die Beziehung zwischen einer Dosis, etwa eines Medikaments, oder einer Stimulusstärke und der Wahrscheinlichkeit einer Reaktion wird gelegentlich durch eine *Weibull-Funktion* beschrieben. Sowohl (A.181) als auch (A.182) können auf die gleiche Weibull-Funktion führen (s. unten).

4. Die Form (A.182) geht für $\beta = 1$ in die Exponentialverteilung $P(Y \leq y) = 1 - \exp(-\lambda y)$ über, mit $\lambda = a^{-1/\beta}$.
5. Die Entscheidung zwischen den Formen (A.181) und (A.182) wird davon abhängen, wie man die zufällige Veränderliche definiert, von deren Werten das jeweils interessierende zufällige Ereignis E abhängt. Will man ausdrücken, dass E eintritt, wenn die zufällige Veränderliche einen Wert *größer* als S , S ein kritischer Wert, ist, wird man auf (A.181) geführt, sofern für die zufällige Veränderliche eine Verteilung vom Weibull-Typ indiziert ist. Macht es Sinn, anzunehmen, dass E eintritt, wenn die zufällige Veränderliche einen Wert kleiner als ein kritischer Wert S annimmt, wird man auf (A.182) geführt, – vorausgesetzt, die zufällige Veränderliche ist vom Weibull-Typ. $\square$

Für die Dichtefunktionen folgen die Ausdrücke

$$f_X(x) = \begin{cases} \frac{\beta_0}{a} \left(\frac{\eta_0 - x}{a}\right)^{\beta_0 - 1} \exp\left[-\left(\frac{\eta_0 - x}{a}\right)^\beta\right], & x \leq \eta_0 \\ 0, & x > \eta_0 \end{cases} \quad (\text{A.183})$$

$$f_Y(y) = \begin{cases} \frac{\beta}{a} \left(\frac{y - \eta_0}{c}\right)^{\beta - 1} \exp\left[-\left(\frac{y - \eta_0}{a}\right)^\beta\right], & y \geq \eta_0 \\ 0, & y < \eta_0 \end{cases} \quad (\text{A.184})$$

und für die Erwartungswerte und Varianzen findet man

$$\mathbb{E}(X) = \eta_0 - \frac{1}{a} \Gamma(1 - 1/\beta), \quad \mathbb{E}(Y) = \eta_0 + \frac{1}{a} \Gamma(1 - 1/\beta) \quad (\text{A.185})$$

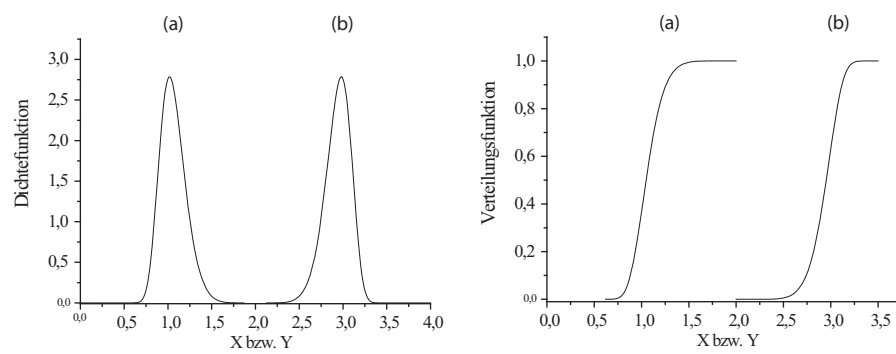
sowie

$$\mathbb{V}(X) = \mathbb{V}(Y) = \frac{1}{a^2} [\Gamma(1 + 2/\beta) - \Gamma^2(1 + 1/\beta)]. \quad (\text{A.186})$$

Verteilungen vom Weibull-Typ ergeben sich auch dann, wenn die Verteilungen der Maxima und Minima unabhängiger zufälliger Größen $X_1, \dots, X_n$ bzw. $Y_1, \dots, Y_n$ gesucht werden; Verteilungen vom Weibull-Typ sind damit auch Extremwertverteilungen, vergl. Abschnitt ??.

Verteilungen vom Weibull-Typ treten nicht nur bei der Verteilung der Bruchfestigkeiten spröden Materials auf, sondern auch, wenn etwa Lebensdauern bzw. Wartezeiten oder Wirkungen der Dosis von Giften bzw. Medikamenten bzw. der Intensität von externen Stimuli auf Sinnesorgane (psychometrische Funktionen) diskutiert werden. Dabei kann es interessant sein, die Verteilungen von Dosiswirkungen und Wartezeiten, oder Stimuluswirkung und Reaktionszeit, zueinander in Beziehung zu setzen, vergl. Abschnitt ??.

Abbildung A.10: Weibulldichten und -verteilungen – (a) nach (A.182), (b) nach (A.181); $\eta_0 = 2$, $a = 1$, $\beta_0 = 7.5$.



Literaturverzeichnis

- [1] Abadi, R.V., Kulikowski, J.J. (1973) Linear summation of spatial harmonics in human vision. *Vision Research* 13, 1625–1628
- [2] Blakemore, C., Campbell, F.W. (1969) On the existence of neurones in the human visual system selectively sensitive to the orientation and size of retinal images. *Journal of Physiology (London)*, 203, 237–260
- [3] Blommaert, F. J. J., Roufs, J. A. J. (1987) Prediction of Thresholds and Latency in the Basis of Experimentally Determined Impulse Responses. *Biological Cybernetics* 56, 329–344
- [4] Campbell, F. W., Gubisch, R.W. (1966) Optical quality of the human eye. *Journal of Physiology*, 186, 558–578
- [5] Campbell, F. W., Kulikowski, J.J., Levinson, J. (1966) The effect of orientation on the visual resolution of gratings. *Journal of Physiology*, 187, 427–436
- [6] Campbell, F. W., Kulikowski, J. J. (1966) Orientational selectivity of the human visual system. *Journal of Physiology*, 187, 437–445
- [7] Campbell, F. W., Carpenter, R. H. S. Levinson, J. Z. (1969) Visibility of aperiodic patterns compared with that of sinusoidal gratings. *Journal of Physiology*, 204, 283 - 298
- [8] Campbell, F. W., Cooper, G. F., Enroth-Cugell, C. (1969) The spatial selectivity of the visual cells of the cat. *The Journal of Physiology (London)*, 203, 223–235
- [9] De Lange Dzn⁵, H. (1952) Experimentss on Flicker and some calculations on an electrical analogue of the foveal systems. *Physica XVIII*, 11, 935–950
- [10] Feller, W.: An Introduction to Probability Theory and its Applications, Vol I. New York, 1968
- [11] Fischer, B. (1972) Optische und neuronale Grundlagen der visuellen Bildübertragung: einheitliche mathematische Behandlung des retinalen Bildes und der Erregbarkeit von Ganglienzellen mit Hilfe der linearen Systemtheorie. *Vision Research*, 12, 1125–1144

⁵Dzn bedeutet Dirks Zoon, also Dirks Sohn. Die Abkürzung Dzn ist also nicht Teil des Namens.

- [12] Graham, N., Nachmias, J. (1971) Detection of grating patterns containing two spatial frequencies: a comparison of single-channel and multiple-channel models. *Vision Research*, 11, 251–259
- [13] Graham, N., Rogowitz, B.E. (1976) Spatial pooling properties deduced from the detectability of FM and quasi-AM gratings: a reanalysis. *Vision Research*, 16, 1021–1026
- [14] Graham, N. (1977) Visual detection of aperiodic spatial stimuli by probability among narrowband spatial frequency channels. *Vision Research*, 17, 637–652
- [15] Graham, N. (1989) *Visual Pattern Analyzers* New York: Oxford University Press (Paperback edition (2001))
- [16] Hauske, G., Wolff, W., Lupp, U. (1976) Matched filters in human vision. *Biological Cybernetics*, 22, 181–188
- [17] Hill, N. J. (2001) Testing hypotheses about psychometric functions. An investigation of some confidence interval methods, their validity, and their use in the assessment of optimal sampling strategies. Thesis submitted in the University of Oxford as a partial requirement of the fulfilment of the degree of Doctor of Philosophy.
- [18] Hines, M. (1976) Line spread function variations near the fovea. *Vision Research*, 16, 567–572
- [19] Hubel, D. H., Wiesel, T. N. (1959) Receptive fields of single neurones in the cat's striate cortex. *Journal of Physiology*, 148, 574–591.
- [20] Hubel, D.N., Wiesel, T.N. (1962). Receptive fields, binocular interaction and functional architecture in the cat's visual cortex. *Journal of Physiology*, 160, 106–154.
- [21] Hubel, D.N., Wiesel, T.N. (1965). Receptive fields and functional architecture in two nonstriate visual areas (18 and 19) of the cat. *Journal of Neurophysiology*, 28, 229–289.
- [22] Kelly, D.H. (1961) Visual Responses to time-dependent stimuli. I. Amplitude sensitivity measurements. *Journal of the Optical Society of America*, 51, 422–429.
- [23] Kelly, D.H. (1971a) Theory of flicker and transient responses 1: uniform fields. *Journal of the Optical Society of America*, 61, 537–546.
- [24] Kelly, D.H. (1971b) Theory of flicker and transient responses 2: counterphase gratings. *Journal of the Optical Society of America*, 61, 632–640
- [25] Kelly, D.H. (1978) Theory of flicker and transient responses. III. An essential nonlinearity. *Journal of the Optical Society of America*, 68, 1481–1490

- [26] Kelly, D.H. (1972) Adaptation effects on spatio-temporal sine-wave thresholds. *Vision Research* 12, 89–101
- [27] Kelly, D.H., Savoie, R.E. (1973) A study of sine-wave contrast sensitivity by two psychophysical methods. *Perception & Psychophysics*, 14, 313–318
- [28] Kelly, D.H. (1975) Spatial frequency selectivity in the retina. *Vision Research*, 15, 665–672
- [29] Kuss, M., Jäkel, F., Wichman, F. A. (2005) Bayesian inference for psychometric functions. *Journal of Vision*, 5, 478–492
- [30] Le Grand, Y.: Light, Colour, and Vision. Chapman and Hall Ttd, London 1968
- [31] Macmillan, N.A., Creelman, N.A.: Detection theory. A user’s guide. Lawrence Erlbaum Associates Ltd, Mahwah 2005
- [32] Maffei, L., Fiorentini, A. (1973) The visual cortex as a spatial frequency analyser. *Vision Research*, 13, 1255–1267
- [33] Manahilov, V., Simpson, W. (1999) Energy model for contrast detection: spatio-temporal characteristics for threshold vision. *Biological Cybernetics*, 81, 61–71
- [34] Marko, H. (1981) the z -Model – a proposal for spatial and temporal modeling of visual threshold perception. *Biological Cybernetics*, 39, 111–123
- [35] Mastronarde, D. N. (1983a) Correlated firing of cat retinal ganglion cells. I. Spontaneously active inputs do X- and Y-cells. *Journal of Neurophysiology*, 49, 303–324
- [36] Mastronarde, D. N. (1983b) Correlated firing of cat retinal ganglion cells. II. Responses to X- and Y-cells to single quantal events. *Journal of Neurophysiology*, 49, 325–349
- [37] Mastronarde, D. N. (1983c) Interactions between ganglion cells in the cat retina. *Journal of Neurophysiology*, 49, 350–356
- [38] Meinhardt, G.: Musterspezifische und stochastische Komponenten bei der Detektion visueller Muster. Aachen 1995
- [39] Mortensen, U.: Wahrscheinlichkeitstheorie. Münster, 1998 (<http://www.uwe-mortensen.de/textauswahl.html>)
- [40] Mortensen, U. (2002) Additive noise, Weibull functions, and the approximation of psychometric functions. *Vision Research*, 42, 2371–2393
- [41] Mostafavi, H., Sakrison, D. J. (1976) Structure and properties of a single channel in the human visual system. *Vision Research*, 9, 957–968
- [42] Nachmias, J., Kocher, E.C. (1970) Visual detection of luminance increments. *Journal of the Optical Societa of America, A*, 60, 382–389

- [43] Nachmias, J. (1981) On the psychometric function for contrast detection. *Vision Research* 21, 215–223
- [44] Nachtigall, C. (1991) Modellierung eines selbstorganisierten Matched-Filters durch ein System gewöhnlicher Differentialgleichungen. Diplomarbeit am Fachbereich Mathematik der Westf. Wilhelms-Universität, Münster
- [45] Peddie, W.: Colour Vision. Edward Arnold, London, 1922
- [46] Quick, R. F. (1974) A vector-magnitude model of contrast detection. *Biological Cybernetics* 16, 65–67
- [47] Rentschler, I., Hilz, R. (1976) Evidence for disinhibition in line detectors. *Vision Research*, 16, 1299–1302
- [48] Röhler, R., Miller, U. Aberl, M. (1965) Zur Messung der Informationsübertragungsfunktion des lebenden menschlichen Auges im reflektierten Licht. *Vision Research*, 9, 407–429
- [49] Roufs, J. A. J. (1972a) Dynamic properties of vision - I: Experimental relationships between flicker and flash thresholds. *Vision Research*, 12, 261–278
- [50] Roufs, J. A. J. (1972b) Dynamic properties of vision - II: Theoretical relationships between flicker and flash thresholds. *Vision Research*, 12, 279–292
- [51] Roufs, J. A. J. (1974) Dynamic properties of vision - VI: Stochastic Threshold Fluctuations and their Relation to Flash-to-Flicker Sensitivity Ratio. *Vision Research*, 14, 871–888
- [52] Sach, M., Nachmias, J., Robson, J. G. (1971) Spatial-Frequency Channels in Human Vision. *Journal of the Optical Society of America*, 61, 1176–1186
- [53] Seelen, W. v. (1968) Informationsverarbeitung in homogenen Netzen von Neuronenmodellen. *Kybernetik*, 5, 133–148
- [54] Shapley, R. M., Tolhurst, D. J. (1973) Edge detectors in human vision. *Journal of Physiology*, 229, 165–183
- [55] Thomas, J. P. (1970) Model of the function of receptive fields in human vision. *Psychological Review*, 77, 121–134.
- [56] Tolhurst, D. J. (1972) On the possibility of edge detector neurones in the human visual system. *Vision Research*, 12, 797–804
- [57] Thomas, R. D. (1999). Assessing sensitivity in a multidimensional space: Some problems and a definition of a general d' . *Psychonomic Bulletin & Review*, 6, 224–238.
- [58] Thomas, R. D. (2003), Further considerations of a general d' in multidimensional space. *Journal of Mathematical Psychology*, 47, 220–224

- [59] Treutwein, B., Strasburger, H. (1999) Fitting the psychometric function. *Perception and Psychophysics*, 61, 87–106
- [60] van Meeteren, A. (1966) A spatial sinewave response of the visual system (a critical literature survey). Soesterberg: Institute for Perception RVO-TNO.
- [61] Watson, A. B. (1981) Derivation of the impulse response: comments on the method of Roufs and Blommaert. *Vision Research* 22, 1335–1337
- [62] Watson, A. B. (1979) Probability summation over time. *Vision Research* 19, 515–522
- [63] Westheimer, G. (1960) Modulation thresholds for sinusoidal light distributions on the retina. *Journal of Physiology*, 152, 67–74
- [64] Whalen, A. D.: Detection of Signals in Noise. Academic Press, New York 1971
- [65] Wichmann, F. A., Hill, N. J. (2001) The psychometric function I: Fitting, sampling and goodness-of-fit. *Perception and Psychophysics*, 63, 1293–1313
- [66] Wichmann, F. A., Hill, N. J. (2001) The psychometric function II. Bootstrap-based confidence intervals and sampling. *Perception and Psychophysics*, 63, 1314–1329
- [67] Wilson, H. R., Giese S. C. (1977) Threshold visibility of frequency grating patterns. *Vision Research*, 17, 177–1190
- [68] Wilson, H. R. (1978) Quantitative characterization of two types of lines-spread function near the fovea. *Vision Research*, 18, 971–981
- [69] Wilson, H.R. (1980) A transducer function for threshold and suprathreshold vision. *Biological Cybernetics* 38, 171–178
- [70] Wilson, H. R., Bergen, F. R. (1979) A four mechanism model for threshold spatial vision. *Vision Research*, 19, 19–32
- [71] Wilson, H. R., Philipps, G. Rentschler, I., Hilz, R. (1979) Spatial probability summation and inhibition in psychophysically measured lines-spread functions. *Vision Research*, 19, 593–598
- [72] Wilson, H. R. (1980) A transducer function for threshold and superthreshold human vision. *Biological Cybernetics*, 88, 171–178

Index

- d' - Sensitivitätsmaß, 90
- Phasenfluß, 6
- Zustandsvariablen, 11

- Aktionspotential, 52
- Aktivierungsrate, 73
- Amplitude, 35
- Amplitudengang, 36
- Approximation
 - lineare, 76
- Argand-Diagramm, 56
- Attenuation, 38, 56
- Attenuationscharakteristik, 112
- Attraktor, 80

- Bahn, 7
- Bifurkation, 80
- Bifurkationspunkt, 80
- Bode-Diagramm, 56

- center frequency, 59

- dB-Einheit, 60
- Dekade, 57, 60
- Differentialgleichung
 - homogene, 31, 46
- Differentiationsoperator, 25
- Differenzier-Glied, 53
- Doppelte Exponentialverteilung, 99
- Dynamik, 6
- Dämpfung
 - schwache, 49
 - starke, 49
- Dämpfungskonstante, 55
 - aktuelle, 55

- Eckfrequenz, 59
- Eigenfrequenz, 23, 26, 49
- Eigenschwingung, 23, 49

- Eigenwert, 78
- Energie, kinetische, 68
- Entscheidungsvariable, 89
- Extremwertverteilung, 99

- Faltung, 31
- Faltungsintegral, 20, 37
- Feedback-Systeme, 41
- Filter, 59
 - Bandpaß-, 59
 - Hochpaß-, 59, 60, 63
 - idealer, 59
 - Schmalband-, 59
 - Tiefpaß-, 59, 63
- Fixpunkt, 7, 8, 11, 73
 - instabiler, 77
 - stabil, 77
- Flimmerverschmelzungsfrequenz, 111
- Flußaxiome, 6
- Flußlinie, 6
- Fokus
 - instabiler, 81
- Fouriertransformierte, 34
- Frequenzfaltung, 180
- Frequenzgang, 36

- Gabor-Funktion, 177
- Gain, 38
- Gamma-Funktion, 50
- Gleichgewichtslage, 11, 73
 - dynamische, 73
- Gleichgewichtszustand, 8
- Grad, 58
- Grenzfall
 - aperiodischer, 26
- Grenzzyklus, 79
 - semistabil, 80
 - stabiler, 80

Heaviside-Funktion, 180
 Hilbertransformierte, 179
 Hintereinanderschaltung, 39
 v. Tief- und Hochpaß, 64
 Hopf-Bifurkation, 81

 Impulsantwort, 31, 36, 59
 Impulsantwortmatrix, 31
 Impulsfunktion, 9
 Induktivität, 53
 instabil, 48
 Integrator, 52
 leckender, 52
 Integrier-Glied, 51

 Jacobi-Matrix, 76

 kanonische Variablen, 48
 kausal, 37
 Kirchhoffsche Regeln, 61
 Koeffizientenmatrix, 11
 Kondensator, 52
 Kontrastinterrelationsfunktion, 142
 Kontrastsensitivität, 138
 Kontrollmatrix, 12
 Korrelationszeit, 192
 Kriechfall, 27

 Laplacetransformierte, 35
 Leistungsspektrum, 71
 Likelihood-Quotient, 85
 line spread function, 110
 Linearisierung, 10, 73, 75
 Linienverwaschungsfunktion, 110
 Lotka-Volterra-Modell, 74

 magnitude function, 97
 Massewirkungsgesetz, 74
 Matrix
 Feedforward, 12
 Impulsantwort, 23
 Modal, 24
 Output, 12
 Mehrfachheit, 47
 Membranpotential, 52
 Minimum-Phasen-Systeme, 45
 Modalmatrix, 24

 Mode, 24
 dynamische, 24

 Normalform, 12
 Normierungskonstanten, 105
 Nullstellen, 46
 Nutzsinalleistung, 90
 Nyquist-Diagramm, 56

 Oktave, 57, 60
 Orbit, 7, 79
 Ortskurve, 56
 Oszillator, harmonischer, 25, 54
 Oszillatorischer Fall, 27

 Parallelschaltung, 40
 Partialbruchzerlegung, 42, 46, 47
 Periode, 7
 Phase, 35
 Phasencharakteristik, 45
 Phasengang, 36, 56
 Phasenkurve, 7
 Phasenporträt
 Beschränkung, 77
 lokales, 77
 Phasenproträt
 globales, 77
 Phasenpunkt, 6
 Phasenraum, 6, 9
 Poincaré – Bendixson, 80
 point spread function, 110, 133
 Poisson-Prozeß, 53
 Polarkoordinaten, 184
 Pole, 46
 Polynom
 n -ten Grades, 47
 charakteristisches, 26, 47, 54
 Potential, 68
 potentielle Energie, 68
 Proportional-Glied, 51
 Prozeß
 autokatalytischer, 74
 psychometrische Funktion, 92, 94
 Punktverwaschungsfunktion, 110, 133

 Quadraturfunktion, 179
 Quicks Modell, 96

- Rauschen
 - bandbegrenztes, 86
 - weißes, 86
- Rauschen, weißes, 87
- Rauschleistung, 90
- RC-Glied, 60
- Receiver-Operator-Characteristics, 85
- Refraktärzeit, 73
- Reihenschaltung, 39
- Relation
 - Eulersche, 37
- relaxieren, 21
- Riccos Gesetz, 132
- ripple-ratio, 112
- Schaltung
 - Hintereinander- oder Reihen-, 25
 - Parallel-, 25
- Sensitivität, 135
- Signal
 - analytisches, 179
- Signal-Rausch-Verhältnis, 90
- spatial probability summation, 100
- Spektralmoment, zweites, 108
- stabil, 48, 59, 77
 - asymptotisch, 77
 - neutral, 77
- Stabilität, 25
- strukturell instabil, 79
- Störung, 11
- Substratinhibition, 74
- Superpositionsexperiment, 137
- Superpositionsprinzip, 20
- System
 - Aktivator-Inhibitor, 73
 - dynamisches, 6
 - homogenes, 78
 - informationsübertragendes, 12
 - kausales, 36
 - konstante Koeffizienten, 11
 - linearer Teil, 76
 - lineares, 10
 - Lotka-Volterra, 78
 - nichtlineares, 10, 73
 - konservative, 68
- Systemfunktion, 36, 46
- Tiefpaß, einpoliger, 59
- Tiefpaßfilter, 53
- Trajektorie, 7
- Transferfunktion, 9, 36, 46, 59
- transiente Effekte, 38
- Vektor
 - Input, 12
 - Output, 12
- Wachstumsrate, 73
- Wahrscheinlichkeitssummation, 92
- Watsons Modell, 103
- Weibull-Verteilung, 194
- Zeitkonstante, 61
- Zentralfrequenz, 59
- Zentrum, 78, 79
- Zustand, 5, 6, 24
- Zustandsfluß, 6
- Zustandsraum, 6, 9
- Zustandstransformation, 6
- Zustandsvektor, 11