

QUANTITATIVE LINGUISTICS

Vol. 5

GLOTTOMETRIKA 3

edited

by

WOLFHART MATTHÄUS



Studienverlag Dr. N. Brockmeyer

Bochum 1980

QUANTITATIVE LINGUISTICS

Editor	G. Altmann, Bochum
Editorial Board	N. D. Andreev, Leningrad M. V. Arapov, Moscow B. Brainerd, Toronto R. Grotjahn, Bochum H. Guiter, Montpellier D. Hérault, Paris E. Hopkins, Bochum W. Lehfeldt, Konstanz W. Matthäus, Bochum R. G. Piotrowski, Leningrad B. Rieger, Aachen/Amsterdam J. Sambor, Warsaw U. Strauss, Bochum D. Wickmann, Aachen

CIP-Kurztitelaufnahme der Deutschen Bibliothek

Glottometrika. – Bochum: Studienverlag Brockmeyer,
3. / Ed. by W. Matthäus. – 1980.
(Quantitative linguistics; Vol. 5)
ISBN 3-88339-136-0
NE: Matthäus, Wolfhart [Hrsg.]

ISBN 3-88339-136-0
Alle Rechte vorbehalten
(c) 1980 by Studienverlag Dr. N. Brockmeyer
Querenburger Höhe 281, 4630 Bochum 1

CONTENTS

GENERAL

ŘEHÁK, J., ŘEHÁKOVÁ, B., Analyse von Kontingenztafeln: Zwei Grundtypen von Aufgaben und das Vorzeichenschema	1
Diskussion: W. Matthäus	29
G. Altmann	29
J. Řehák, B. Řeháková	31

PHONOLOGY

GROTJAHN, R., Einige kritische Bemerkungen zu W. Lehfeldt "Zur numerischen Erfassung der Schwierigkeit des Sprechbewegungsablaufs"	34
--	----

TEXT ANALYSIS

BORODA, M.G., Häufigkeitsstrukturen musikalischer Texte	36
Diskussion: G. Altmann	70
E. Fischer	71
W.-D. Schäfer	73
J. Wildgruber	78
M.G. Boroda	79

PSYCHOLINGUISTICS

MARX, W., STRUBE, G., The Butterfly Revisited. A re-analysis of Deese's study 'On the structure of associative meaning'	97
KOLMAN, L., Gedächtnis und Erkennen. Ein Beitrag zur semantischen Theorie der Erkennungsprozesse	105

BIBLIOGRAPHY

STOFFER, T.H., BORODA, M.G., Mathematisierungstendenzen in der Musikwissenschaft: Eine Bibliographie zur quantitativ-statistischen und algebraisch-formalen Analyse musikalischer Strukturen	198
--	-----

CURRENT BIBLIOGRAPHY	223
----------------------	-----

LIST OF CONTRIBUTORS

INSTRUCTIONS TO AUTHORS

We would like to express our gratitude to the

STIFTUNG VOLKSWAGENWERK

a generous grant from whom made possible the translation of Řehák & Řeháková 'Analyse von Kontingenztafeln: Zwei Grundtypen von Aufgaben und das Vorzeichenschema' and Boroda 'Häufigkeitsstrukturen musikalischer Texte', as well as to the

DEUTSCHE FORSCHUNGSGEMEINSCHAFT

whose generous support permitted the translation of Kolman, 'Gedächtnis und Erkennen: Ein Beitrag zur semantischen Theorie der Erkennungsprozesse' to be carried out.

ANALYSE VON KONTINGENZTAFELN: ZWEI GRUNDTYPEN VON AUFGABEN UND DAS VORZEICHENSHEMA

J. Řehák, B. Řeháková, Prag

Die Kontingenztafel gehört zur methodischen Grundausstattung jedes mit Daten arbeitenden Sozialwissenschaftlers. Sie ist auch die üblichste Anordnungsweise, in der wir Daten vom Rechner bekommen und in Forschungsberichten darstellen. Über Kontingenztafeln bzw. über die Auswertung der Doppelklassifikation wurde schon viel geschrieben, nichtsdestoweniger machen wir bei ihrer Verarbeitung oft Fehler. Es ist das Ziel dieser Arbeit, auf zwei grundlegende Tafeltypen, vor allem aber auf das Zeichenschema aufmerksam zu machen.

Das Zeichenschema in den Kontingenztafeln ist ein sehr geeignetes Mittel für die Analyse der Beziehungstypen in der Kontingenztafel und gehört zu den praktischsten Neuerungen unserer soziologischen Methodologie. Es wurde von Linhart und Šafár (1967) eingeführt und beruht auf dem intuitiven Testen der Übereinstimmung zwischen der beobachteten und der erwarteten Häufigkeit in jeder Zelle der Kontingenztafel. Aufgrund dieses Testens druckt der Rechner ein Schema der folgenden Zeichen aus: +++, ++, +, 0, -, --, ---. Das Zeichen "+" bedeutet, daß die beobachteten Häufigkeiten die erwarteten Häufigkeiten signifikant übersteigen, das Zeichen "-" bedeutet, daß die beobachteten Häufigkeiten signifikant kleiner sind als die erwarteten, und 0 bedeutet einen nichtsignifikanten Unterschied. Die Anzahl der Zeichen wird nach der Signifikanz des Unterschiedes abgestuft, z.B. bedeutet "+" eine Signifikanz auf dem Niveau 0.05, "++" auf dem Niveau 0.01 und "+++" auf dem Niveau 0.001 (analog "-"). Die Wahl der Signifikanzgrenzen steht im Belieben des Benutzers (oder des Programmierers), ist aber durch eine gewisse Konvention und Gewohnheit geregelt. Die Autoren des Schemas (Linhart/Šafár 1967) haben für das Testen eine intuitive Statistik eingeführt, von der sie annehmen, daß sie asymptotisch normal verteilt ist und daher die bekannten Eigenschaften des sogenannten z-score hat.

Bezeichnen wir mit n_{ij} die absoluten Häufigkeiten in den Zellen einer $R \times C$ -Tabelle, mit $n_{i.}$ die Zeilensummen, mit $n_{.j}$ die Spaltensummen und mit n die gesamte Anzahl der Beobachtungen.¹⁾ Linhart und Šafář schlagen die folgende Statistik vor:

$$z_{ij} = \frac{n_{ij} - \frac{n_{i.} \cdot n_{.j}}{n}}{\sqrt{\frac{n_{i.} \cdot n_{.j}}{n} \left(1 - \frac{n_{i.} \cdot n_{.j}}{n^2}\right)}} \quad (1)$$

$$= \sqrt{n} \frac{nn_{ij} - n_{i.} \cdot n_{.j}}{\sqrt{n_{i.} \cdot n_{.j} (n^2 - n_{i.} \cdot n_{.j})}} \quad (2)$$

Die graphische Übersicht hat dank ihrer Vorteile - schnelle Orientierung und bequeme Beurteilung der qualitativen Beziehungsstruktur zweier Variablen - breite Verwendung gefunden. Die Idee hat sich schnell in verschiedenen Programmen mit diversen Modifikationen und Fehlern niedergeschlagen. Das Ziel dieser Arbeit ist es, richtige Formeln für diese Methode einzuführen und sie dadurch zu revidieren und zu vereinheitlichen. Die Resultate folgen aus der statistischen Theorie asymptotischer Tests, d.h. Tests für große Stichproben (vgl. Rao 1973, bes. Kap. 6). Es werden hier die zwei üblichsten Typen der Kontingenztabellen analysiert. Es wird gezeigt, daß für die beiden diskutierten Aufgabentypen die formale Berechnung des Vorzeichenschemas und der Testcharakteristika gleich ist, obwohl die Tests der Grundhypothesen und das Vorzeichenschema unterschiedliche Bedeutungen haben. Die mathematische Argumentation und die technischen Aspekte sind im letzten Teil des Artikels und in einigen Anmerkungen enthalten, die an Spezialisten gerichtet sind, während der Hauptteil des Texts der praktischen Arbeit bei der Datenanalyse gewidmet ist.

Wir führen noch einige weitere Bezeichnungen ein:

$$f_{j/i} = \frac{n_{ij}}{n_{i.}} = \text{zeilenweise relativierte Zellenhäufigkeit}$$

$$f_{i.} = \frac{n_{i.}}{n} = \text{relative Häufigkeit der Zeilenvariablen in der Stichprobe (Randverteilung der Zeilensummen)}$$

$$f_{.j} = \frac{n_{.j}}{n} = \text{relative Häufigkeit der Spaltenvariablen in der Stichprobe (Randverteilung der Spaltensummen)}$$

$$f_{ij} = \frac{n_{ij}}{n} = \text{relative Zellenhäufigkeit in Bezug auf den gesamten Stichprobenumfang.}$$

$p_{j/i}$, $p_{i.}$, $p_{.j}$, p_{ij} haben die analoge Bedeutung für die Verteilungen in der Gesamtheit, während die f den Stichprobenwerten entsprechen.

1. VERGLEICH MEHRERER STICHPROBEN - TYP A

Die in einer Kontingenztafel geordneten Daten repräsentieren unabhängige Gesamtheiten, die wir in Bezug auf ein nominales Merkmal vergleichen wollen. Seien R solche Gesamtheiten gegeben, und das nominale Merkmal habe C Werte. Jede Zeile entspricht einer Gesamtheit und enthält die Häufigkeitsverteilung, die durch eine einfache zufällige Stichprobenerhebung²⁾ aus der gegebenen Gesamtheit entsteht.

BEISPIELE:

- (a) Gesamtheiten: Länder, Arbeiter ausgewählter Betriebe;
Merkmal: Typ der Beziehung des Arbeiters zur Arbeit.
- (b) Gesamtheiten: Bezirke, erwachsene Einwohner;
Merkmal: Typ der Leseraktivität in Bezug zu Wochenblättern.
- (c) Gesamtheiten: Jugend in Prag, die restliche Jugend in der ČSSR;
Merkmal: Typ der Freizeitgestaltung.

Dieser Fall ist sehr häufig. Seine Charakterisierung besteht darin, daß

- (a) die Daten in den Zeilen unabhängig entstehen;
- (b) die Zeilensumme durch den Stichprobenumfang in der gegebenen Gesamtheit im voraus bestimmt wird;
- (c) die Größe der Stichproben zu den Größen der Grundgesamtheiten nicht proportional zu sein braucht;
- (d) die gemeinsame Verteilung der Häufigkeiten nicht spezifiziert wird.

Hypothese: Alle Grundgesamtheiten haben dieselbe Verteilung, die Unterschiede in den Stichproben kann man durch die Zufälligkeit der Stichprobe oder durch zufällige Einflüsse bei der Genese einzelner Daten (Meßfehler oder Prozeß des Entstehens der Stichprobe) erklären.

Formal kann man die Hypothese formulieren als

$$p_{j/i} = p_{.j} \text{ für alle } i \text{ und } j \quad (3)$$

(alle R durch die Tabelle repräsentierten Grundgesamtheiten haben dieselbe Verteilung); die Zahlen $p_{.j}$ sind nicht spezifiziert.

Alternativhypothese: Mindestens eine Gesamtheit unterscheidet sich in ihrer Verteilung von den anderen, und zwar bei der relativen Häufigkeit mindestens eines Wertes des Merkmals; mindestens eine Zeile stellt eine von den anderen unterschiedliche Häufigkeitsverteilung der Grundgesamtheit dar.

Die Alternativhypothese umfaßt also alle Fälle, die sich von dem Fall der Gleichheit der Verteilungen unterscheiden. Es handelt sich also um eine omnibus (allgemeine) Alternativhypothese.

BEISPIEL 1.

Vergleichstabelle der Verteilungen von Typen der Freizeitgestaltung der Jugend (Typen A, B, C, D) für zwei Stichproben: I. Jugend der ČSSR (außer Prag), II. Jugend aus Prag.

Tabelle 1. Absolute Häufigkeiten und Zeilenproportionen (in Klammern)³

	Typ				Umfang	Anteil der Stichprobe
	A	B	C	D		
Stichprobe 1	236 (20.8%)	370 (32.7%)	384 (33.9%)	143 (12.6%)	1133	56.8%
Stichprobe 2	159 (18.4%)	320 (37.1%)	232 (26.9%)	151 (17.5%)	862	43.2%
Insgesamt	395 (19.8%)	690 (34.6%)	616 (30.9%)	294 (14.7%)	1995	

2. TEST DER UNABHÄNGIGKEIT ZWEIER NOMINALER MERKMALE - TYP B

Die Daten in der Kontingenztafel stammen aus einer einfachen zufälligen Stichprobe, die Zeilen und die Spalten entsprechen zwei unterschiedlichen nominalen Merkmalen, deren statistische Unabhängigkeit getestet werden soll. Dieser Fall kann noch analytisch weiter spezifiziert werden, inwiefern uns entweder die gegenseitige statistische Abhängigkeit (symmetrische statistische Beziehung) oder die asymmetrische Abhängigkeit eines Merkmals von dem anderen interessiert. Im Unterschied zur Bestimmung der Koeffizienten der Zusammenhänge, die in beiden Fällen unterschiedlich gerechnet werden (vgl. z.B. Řehák/Řeháková 1973), führt man die Unabhängigkeitstests auf dieselbe Weise durch. Die Unterscheidung beider Fälle liegt in der Interpretation; sie ist empirischer, nicht statistischer Natur. Beispiele sind aus der laufenden Forschungspraxis bekannt.

BEISPIELE:

- (a) Merkmale: Geschlecht, Typ der Freizeitgestaltung;
Gesamtheit: Jugend der ČSSR;
- (b) Merkmale: Form der Beteiligung an der Betriebsleitung, in-

haltliche Ausrichtung der Beteiligung an der Betriebsleitung;
Gesamtheit: Arbeiter der Nahrungsmittelindustrie.

- (c) Merkmale: Typ der Arbeitsorganisation, Typ der interpersonellen Beziehungen; Gesamtheit: Forschungskollektive.

Charakterisierung des Typs B:

- (a) Die Daten der ganzen Tabelle entstehen als die Realisierung einer zufälligen Stichprobe⁴⁾ aus einer Gesamtheit.
 (b) Die Gesamtsumme n wird vor der Erhebung festgelegt, die Zeilen- und Spaltensummen dagegen nicht.
 (c) Man erwartet, daß die absoluten Zeilen- und Spaltensummen, im Rahmen der statistischen Toleranz, der Situation in der Grundgesamtheit proportional entsprechen.
 (d) Die Zeilen- und die Spaltenrandverteilung ist nicht spezifiziert.

Hypothese: Die beiden Merkmale in der gegebenen Gesamtheit hängen nicht zusammen, daher entspricht ihre empirische Häufigkeit der statistischen Unabhängigkeit. Als Formel kann man diese Hypothese folgendermaßen formulieren:

$$p_{ij} = p_{i.} \times p_{.j} \text{ für alle } i, j \quad (4)$$

$p_{i.}$, $p_{.j}$ sind nicht spezifiziert.

Die alternative Omnibushypothese: Mindestens bei einem Paar (i, j) ist die statistische Unabhängigkeit gestört,

$$p_{ij} \neq p_{i.} \times p_{.j} \quad (5)$$

Sie enthält also alle möglichen Fälle der unterschiedlichsten Abhängigkeitstypen.

BEISPIEL 2:

Für die Gesamtheit der Jugend der ČSSR (außer Prag) untersuchen wir den Zusammenhang des Vorkommens zweier Typologien (Freizeitgestaltung: Typen A, B, C, D, und Interessenorientierung: Typen a, b, c). Die gemeinsame Verteilung der absoluten und der relativen

Häufigkeiten ist in der Tabelle 2 angegeben.

Tabelle 2. Gemeinsame absolute und relative Verteilung von zwei Merkmalen⁵⁾

		Typ				Insgesamt
		A	B	C	D	
Typ	a	100 (9.1%)	32 (2.9%)	151 (13.8%)	124 (11.3%)	407 (37.1%)
	b	51 (4.6%)	16 (1.5%)	103 (9.4%)	40 (3.6%)	210 (19.1%)
	c	118 (10.8%)	73 (6.7%)	159 (14.5%)	130 (11.9%)	480 (43.8%)
Insgesamt		269 (24.5%)	121 (11.0%)	413 (37.6%)	294 (26.8%)	1097 (100%)

3. TEST DER NULL-HYPOTHESE FÜR BEIDE FÄLLE (TYP A UND B)

Auch wenn sich die Aufgaben A und B durch ihr analytisches Ziel und durch ihre Einbettung in den soziologischen Kontext der Dateninterpretation unterscheiden, führt man den Test der beiden Null-Hypothesen formal identisch durch, indem man die bekannte⁶⁾ Chi-Quadrat (χ^2) Statistik berechnet. Zur Übersicht und auch für praktische Zwecke bringen wir eine Reihe von äquivalenten Formeln⁷⁾:

$$\chi^2 = \sum \sum \frac{\left(n_{ij} - \frac{n_{i.} \cdot n_{.j}}{n} \right)^2}{\frac{n_{i.} \cdot n_{.j}}{n}} \quad (6)$$

$$= \frac{1}{n} \sum \sum \frac{(n n_{ij} - n_{i.} \cdot n_{.j})^2}{n_{i.} \cdot n_{.j}} \quad (7)$$

$$= n \left(\sum \frac{n_{ij}^2}{n_{i.} \cdot n_{.j}} - 1 \right) \quad (8)$$

$$= n \left(\sum \frac{n_{ij} \cdot f_{j/i}}{n_{.j}} - 1 \right) \quad (9)$$

$$= n \left(\sum \frac{f_{ij}^2}{f_{i.} \cdot f_{.j}} - 1 \right) \quad (10)$$

Die Anzahl der Freiheitsgrade $FG = (R-1)(C-1)$. (11)

Das Verfahren:

1. Man bestimmt die Null-Hypothese H_0 und das Signifikanzniveau α .
2. Man berechnet X^2 und bestimmt die FG.
3. Man findet den kritischen Wert $X_{\alpha;FG}^2$ in statistischen Tabellen.
4. Man entscheidet:

Wenn $X^2 \geq X_{\alpha;FG}^2$, dann lehnt man H_0 ab. (12)

Andernfalls hat man keinen Grund zur Ablehnung. Alternativ kann man auch so verfahren, daß man zu X^2 und zu FG das Signifikanzniveau α^* findet, und je nach seiner Höhe entscheidet, ob man die Hypothese annimmt oder ablehnt, d.h. wenn $\alpha^* \leq \alpha$ (im voraus gewählt), dann lehnt man die Hypothese ab. Dieses Verfahren ist in der letzten Zeit beim Rechneroutput üblich geworden (es ist einfacher, es zu programmieren, und der output dient allen Benutzern ohne Bezug auf ihre Wahl von α).

Einschränkungen bei der Anwendung dieses Tests:⁸⁾

1. Kein erwarteter Wert dürfte kleiner als 1 sein; wenn ein Drittel oder ein Viertel der erwarteten Werte zwischen 1 bis 5 liegt, so kann man den Test anwenden (Lancaster 1969).
2. $2 \times C$ (bzw. $R \times 2$)-Tafeln kann man mit diesem Test testen, wenn alle erwarteten Werte größer oder gleich 1 sind; auch diese

Regel ist konservativ, man kann die Grenze bis auf 0,5 reduzieren (Lewontin/Felsenstein 1965).

BEISPIEL 2 (Fortsetzung):

Die X^2 -Statistik, berechnet nach der Formel (8), ergibt⁹⁾

$$X^2 = 1097 \left[\frac{100^2}{407(269)} + \frac{32^2}{407(121)} + \dots + \frac{130^2}{480(294)} - 1 \right]$$

$$= 1097(1.027507345-1)$$

$$= 30.18$$

$$R = 3, C = 4, FG = (R-1)(C-1) = 6$$

Die Signifikanz kann man in den Tabellen der χ^2 -Verteilung feststellen (vgl. z.B. Janko 1958, Tab. 4,5). X^2 ist hoch signifikant, α ist kleiner als 0.0001.

4. DAS ZEICHENSCHEMA¹⁰⁾

Wenn die Null-Hypothese abgelehnt wird, dann interessiert uns die Struktur der Abhängigkeit (Typ B) oder des Unterschiedes der Gesamtheiten (Typ A). Zu diesem Zweck kann man die Abweichungen der beobachteten von den erwarteten Werten in den einzelnen Feldern der Tafel testen. Zwar benutzt man dabei das Verfahren des Hypothesentestens, aber es handelt sich nicht wirklich um den Test einer spezifischen Hypothese. Vielmehr geht es darum, wie die akzeptierte globale Alternativhypothese der Existenz eines Unterschiedes oder einer Abhängigkeit zu spezifizieren ist. Das folgende Verfahren ist daher explorativ.

Obwohl die beiden Typen A und B unterschiedlichen empirischen als auch statistischen Situationen entsprechen, wird das Zeichen-

schema (ebenso wie der Test) in beiden Fällen formal identisch bestimmt. Die Ableitung erfolgt jedoch aus zwei unterschiedlichen Modellen und wird, ebenso wie die Ableitung des χ^2 -Tests, unterschiedlich durchgeführt (vgl. §6).

Das Verfahren besteht aus folgenden Schritten:

1. Man bestimmt drei Signifikanzebenen (z.B. 0.05; 0.01; 0.001 oder 0.1; 0.05; 0.01) und findet die dazugehörigen kritischen Werte der Normalverteilung (je nach Vereinbarung entweder die Werte für den einseitigen oder für den zweiseitigen Test; wir empfehlen den zweiseitigen Test)¹¹⁾.
2. Man berechnet den z-score für jedes Feld (i,j) der Kontingenztafel nach einer der folgenden Formeln

$$z_{ij} = \frac{N_{ij} - n f_{i.} f_{.j}}{\sqrt{n f_{i.} (1 - f_{i.}) f_{.j} (1 - f_{.j})}} \quad (13)$$

$$= n \frac{n_{ij} - \frac{n_{i.} n_{.j}}{n}}{\sqrt{\frac{n_{i.} n_{.j}}{n} (n - n_{.j}) (n - n_{i.})}} \quad (14)$$

$$= \sqrt{n} \frac{n n_{ij} - n_{i.} n_{.j}}{\sqrt{n_{i.} n_{.j} (n - n_{.j}) (n - n_{i.})}} \quad (15)$$

$$= \sqrt{n} \frac{n_{ij} (n - n_{i.} - n_{.j} + n_{ij}) - (n_{i.} - n_{ij}) (n_{.j} - n_{ij})}{\sqrt{n_{i.} n_{.j} (n - n_{.j}) (n - n_{i.})}} \quad (16)$$

3. Die Resultate druckt man¹²⁾ nach der in der Tabelle A angegebenen Regel.

Tabelle A

Drucke	wenn z_{ij} in diesem Intervall liegt
+++	$< z_3, \infty$
++	$< z_2, z_3$
+	$< z_1, z_2$
0	$(-z_1, z_1)$
-	$(-z_2, -z_1)$
--	$(-z_3, -z_2)$
---	$(-\infty, -z_3)$

Dieses Verfahren hat allerdings nur dann Sinn, wenn wir die Hypothese des Unterschiedes der Gesamtheiten bzw. die Hypothese der Abhängigkeit der Merkmale angenommen haben und spezifizieren wollen, wo diese Unterschiede bzw. Abhängigkeiten in Erscheinung treten.

Bemerkungen zu der Methode

1. Der Test gilt für jedes Feld separat. Das Schema als ganzes, d.h. die Struktur aller Testresultate aus einzelnen Feldern, hat eine viel niedrigere statistische Glaubwürdigkeit als jedes einzelne Feld, da wir die gleichzeitige Gültigkeit aller Einzelresultate, die alle mit Fehlern belastet sind, voraussetzen. Daher kann man das Zeichenschema nicht als ein voll aussagekräftiges, sondern nur als ein graphisches Hilfsmittel für die analytische Arbeit betrachten, und man muß immer beachten, daß diese Resultate mit größeren Fehlerwahrscheinlichkeiten belastet sind, als man gewöhnt ist.

2. Das Verfahren ist mit dem χ^2 -Test mit 1 FG identisch, wenn man ihn auf eine Vierfeldertafel anwendet; die Tafel konstruiert man aus der ursprünglichen Tafel so, daß das gegebene Feld ein Feld bildet und die restlichen durch Vereinigung der übrigen Zeilen und Spalten entstehen. Das Zeichen muß aber gesondert bestimmt werden.

BEISPIEL 2 (Fortsetzung):

Wir bestimmen das Zeichenschema für die Tafel der Abhängigkeit zweier Typologien aus dem Beispiel 2. Wir berechnen die z-Werte für einzelne Zellen nach der Formel (15). Beispielsweise bekommen wir für z_{11} :¹³⁾

$$z_{11} = \sqrt{\frac{1097}{407(1097-407)}} \cdot \frac{[1097(100)-407(269)]}{\sqrt{269(1097-269)}} = 0.03.$$

Die Tabelle der z-Werte für die Tabelle 2 ist

	A	B	C	D
a	0.03	-2.57	-0.29	2.11
b	-0.09	-1.75	+3.79	-2.82
c	0.04	3.90	-2.73	0.19

Die Tabelle der Zeichen für $z_1 = 1.96$, $z_2 = 2.58$, $z_3 = 3.29$ ($\alpha_1 = 0.05$, $\alpha_2 = 0.01$, $\alpha_3 = 0.001$):

	A	B	C	D
a	0	--	0	+
b	0	0	+++	--
c	0	+++	--	0

Typ A trägt zum Zusammenhang nicht bei. Korreliert sind miteinander: (b,C), (c,B) und (a,D), die häufiger als erwartet auftreten, sowie (a,B), (b,D) und (c,C), die seltener als erwartet auftreten.

3. Die Beschränkung auf die Berechnung der z-Werte folgt aus der Approximation der Binomialverteilung durch die Normalverteilung.

Um diese Approximation benutzen zu können, sollte die erwartete Häufigkeit (d.h. $n_{i.}p_{.j}$ für den Fall A und $np_{i.}p_{.j}$ für den Fall B) größer oder gleich 5 sein. Bei kleineren Werten sollte

man den exakten Binomialtest berechnen und sein Signifikanzniveau bestimmen. In der Praxis ersetzt man die unbekannten Parameter $p_{i.}$, $p_{.j}$ durch ihre Schätzungen aus der Stichprobe, $f_{i.}$, $f_{.j}$. Im Falle A testet man dann die Hypothese, daß die binomialverteilte Gesamtheit mit dem Umfang $n_{i.}$ der Wahrscheinlichkeit $P = p_{.j}$ entspricht. Im Falle B testet man die Hypothese, daß die binomialverteilte Gesamtheit mit dem Umfang n der Wahrscheinlichkeit $P = p_{i.}p_{.j}$ entspricht. Wir berechnen einseitige oder zweiseitige Tests je nachdem, wie wir die Signifikanzniveaus bei der Normalapproximation gewählt haben, damit die α -Werte mit einseitigen oder zweiseitigen Grenzen der z-Werte vergleichbar sind.

5. VEREINFACHUNG DER RECHNUNG BEI 2x2 (bzw. Rx2)-TAFELN

Für die 2x2 (bzw. nach der Transposition die Rx2)-Tafel kann man die Berechnung des χ^2 vereinfacht durchführen. Dieser Fall ist häufig, er entspricht dem Vergleich zweier multinomial verteilter Gesamtheiten und dem Vergleich von R Binomialverteilungen. Außer Formeln (6) bis (10) kann man eine der folgenden verwenden. Berechnen wir (für die 2x2-Tafel)

$$g_j = \frac{n_{1j}}{n_{.j}}, \quad g = \frac{n_{1.}}{n},$$

dann wird

$$\chi^2 = \frac{1}{g(1-g)} \sum n_{.j} (g_{j.} - g)^2 \quad (17)$$

$$= \frac{1}{g(1-g)} \left(\sum n_{.j} g_{j.}^2 - ng^2 \right) \quad (18)$$

$$= \frac{1}{g(1-g)} \left(\sum n_{1j} g_j - n_{1.} g \right). \quad (19)$$

Gleichzeitig vereinfacht sich auch die Berechnung des Zeichenschemas: Es reicht, Zeichen für die erste Zeile zu berechnen, die z-Werte der zweiten Zeile sind dieselben Zahlen mit umgekehrtem Vorzeichen, d.h.

$$z_{1j} = -z_{2j}. \quad (20)$$

BEISPIEL 1 (Fortsetzung):

Zur Berechnung des χ^2 benutzen wir die Formel (19):

$$\begin{aligned} \chi^2 &= \frac{1}{g(1-g)} \left(\sum \frac{n_{1j}^2}{n_{.j}} - \frac{n_{1.}^2}{n} \right) = \\ &= \frac{1995^2}{1133(862)} \left(\frac{236^2}{395} + \frac{370^2}{690} + \frac{384^2}{616} + \frac{143^2}{294} - \frac{1133^2}{1995} \right) = \\ &= 19.967 \end{aligned}$$

mit $FG = 3$.

In den Tabellen der χ^2 -Verteilung findet man, daß $\alpha^* = 0.003$, d.h. χ^2 ist hoch signifikant, wir nehmen die Alternativhypothese des Unterschiedes beider Gesamtheiten an. Das Zeichenschema zeigt, bei welchen Typen sich die beiden Gesamtheiten signifikant unterscheiden. Es reicht, die z-Werte der ersten Zeile zu bestimmen (die zweite enthält dieselben Zahlen mit umgekehrtem Vorzeichen):

	A	B	C	D
Gesamtheit I	1.32	-2.08	3.34	-3.06
Gesamtheit II	-1.32	2.08	-3.34	3.06

Das Zeichenschema :

	A	B	C	D
Gesamtheit I	0	-	+++	---
Gesamtheit II	0	+	---	+++

Der größte Unterschied besteht bei den Typen C und D, ein kleinerer, jedoch auch signifikanter beim Typ B.

Noch mehr vereinfacht sich die Berechnung des χ^2 und des Zeichenschemas bei einer Vierfeldertafel.

$$\chi^2 = n \frac{(n_{11}n_{22} - n_{12}n_{21})^2}{n_{1.}n_{2.}n_{.1}n_{.2}} \quad (21)$$

$$\text{mit } FG = 1 \quad (22)$$

$$z_{11} = \sqrt{n} \frac{n_{11}n_{22} - n_{12}n_{21}}{\sqrt{n_{1.}n_{2.}n_{.1}n_{.2}}} \quad (23)$$

$$z_{11} = z_{22} = -z_{12} = -z_{21} \quad (24)$$

$$z_{ij}^2 = \chi^2 \quad (25)$$

Die Signifikanz des χ^2 bestimmt also das Zeichenschema, die beiden Verfahren sind äquivalent. Es reicht, ein einziges Zeichen zu bestimmen, die anderen folgen automatisch.

6. ABLEITUNG DER FORMEL FÜR DEN TEST DER ABWEICHUNG IN EINER EINZELNEN ZELLE.

Ohne Einschränkung der Allgemeinheit reicht es, die Ableitung für die Zelle (1,1) zu zeigen. Man muß aber die Modelle A und B unterscheiden.

Typ A

Nehmen wir an, daß die Verteilungen in den Zeilen unabhängig sind, daß die Nullhypothese $p_{j/i} = p_{.j}$ gilt, daß $p_{.j}$ unbekannt ist und $p_{i.} = \frac{n_{i.}}{n}$ im voraus gegeben ist. Man kann nach zwei Methoden verfahren:

Methode 1:

Wir vergleichen die zeilenbezogene relative Häufigkeit in einer Zelle mit der relativen Häufigkeit des Restes der Gesamtheit, d.h. wenn z.B. $n_{11}/n_{1.}$ die relative Häufigkeit in der Zelle (1,1) ist, dann ist die absolute Häufigkeit der gegebenen Kategorie im Rest der Stichprobe $n_{.1} - n_{11}$; der Umfang des dazugehörigen Komplements ist dann $n - n_{1.}$ und die zu vergleichende relative Häufigkeit ist $(n_{.1} - n_{11}) / (n - n_{1.})$. Der Test der Abweichung in einer Zelle beruht auf dem Vergleich dieser beiden relativen Häufigkeiten, d.h.

$$w = \frac{n_{11}}{n_{1.}} - \frac{n_{.1} - n_{11}}{n - n_{1.}} \quad (26)$$

Beide Größen in (26) sind asymptotisch normal verteilt (da sie beide aus einer Multinomialverteilung stammen, die man für große n mit der Normalverteilung approximieren kann), daher ist auch ihre Differenz w asymptotisch normal verteilt. Die Varianz beider Größen und daher auch die ihres Unterschiedes berechnet man aus der ursprünglichen Multinomialverteilung. Es ist

$$\text{Var } w = \text{Var } \frac{n_{11}}{n_{1.}} - 2 \text{Cov} \left(\frac{n_{11}}{n_{1.}}, \frac{n_{.1} - n_{11}}{n - n_{1.}} \right) + \text{Var} \left(\frac{n_{.1} - n_{11}}{n - n_{1.}} \right).$$

Da wir zwei unabhängige Häufigkeiten vergleichen und $n_{1.}$ und $n - n_{1.}$ festgelegt sind, verschwindet die Kovarianz. Wegen

$$E \left(\frac{n_{11}}{n_{1.}} \right) = E \left(\frac{n_{.1} - n_{11}}{n - n_{1.}} \right) = p_{.1}$$

(woraus $E(w) = 0$ folgt)

und wegen

$$\text{Var} \left(\frac{n_{11}}{n_{1.}} \right) = \frac{1}{n_{1.}} p_{.1} (1 - p_{.1})$$

$$\text{Var} \left(\frac{n_{.1} - n_{11}}{n - n_{1.}} \right) = \frac{1}{n - n_{1.}} p_{.1} (1 - p_{.1})$$

bekommt man mit etwas Algebra

$$\begin{aligned} \text{Var } w &= p_{.1} (1 - p_{.1}) \left[\frac{1}{n_{1.}} + \frac{1}{n - n_{1.}} \right] \\ &= \frac{1}{n} \frac{p_{.1} (1 - p_{.1})}{p_{1.} (1 - p_{1.})} \end{aligned} \quad (27)$$

Die Größe $z = \frac{w}{\sqrt{\text{Var } w}}$ ist also normal verteilt mit der Erwartung 0 und Standardabweichung 1. Die Schätzung von $\text{Var } w$ erhält man dadurch, daß man für die Wahrscheinlichkeit $p_{.1}$ die empirische relative Häufigkeit $f_{.1}$ einsetzt:

$$w = \frac{n_{11}}{n_{1.}} - \frac{n_{.1} - n_{11}}{n - n_{1.}} = \frac{n_{11} (n - n_{1.}) - n_{1.} (n_{.1} - n_{11})}{n_{1.} (n - n_{1.})} = \frac{n \cdot n_{11} - n_{1.} n_{.1}}{n_{1.} (n - n_{1.})}$$

$$\text{Var } w = \frac{1}{n} \cdot \frac{n_{.1}(n-n_{.1})}{n_{1.}(n-n_{1.})}$$

$$z = \frac{w}{\sqrt{\text{Var } w}} = \frac{(n \cdot n_{11} - n_{1.} \cdot n_{.1}) \cdot \sqrt{n} \cdot \sqrt{n_{1.}(n-n_{1.})}}{n_{1.}(n-n_{1.}) \cdot \sqrt{n_{.1}(n-n_{.1})}}$$

und schließlich

$$z = \sqrt{n} \frac{n \cdot n_{11} - n_{1.} \cdot n_{.1}}{\sqrt{n_{1.} \cdot n_{.1} (n-n_{1.}) (n-n_{.1})}} \quad (28)$$

was mit der Formel (15) identisch ist.

Methode 2:

Das zweite Verfahren beruht auf dem Vergleich der Häufigkeit des gegebenen Feldes mit der relativen Häufigkeit der betreffenden Spalte. Während wir im vorigen Fall einen Teil mit seinem Komplement verglichen haben, kontrastieren wir hier einen Teil mit dem Ganzen. Für diesen Zweck benutzen wir die Differenz

$$w' = \frac{n_{11}}{n_{1.}} - \frac{n_{.1}}{n} \quad (29)$$

Wir verfahren hier genauso wie bei (27). Man muß hier jedoch beachten, daß die Größen nicht unabhängig sind, so daß die Kovarianz nicht verschwindet. Deswegen modifizieren wir w' erst so, daß die beiden Größen unabhängig werden, dann leiten wir die Varianzen ab und stellen die normalisierte Teststatistik auf. Wir schreiben (29) als

$$w' = \frac{n_{11}}{n_{1.}} - \frac{n_{.1}}{n}$$

$$= \frac{n_{11}}{n_{1.}} - \frac{n_{11} + n_{.1} - n_{11}}{n}$$

$$= \left(\frac{n_{11}}{n_{1.}} - \frac{n_{11}}{n} \right) - \left(\frac{n_{.1} - n_{11}}{n} \right)$$

Die erste Klammer enthält als Zufallsgröße nur die Häufigkeit n_{11} , die zweite Klammer als Zufallsgröße nur den Rest (das Komplement) der Spaltenhäufigkeit $n_{.1}$. Diese beiden komplementären Zufallsgrößen gehören zu verschiedenen Stichproben und sind daher unabhängig, ihre Kovarianz ist gleich 0. Wir haben also

$$\text{Var} \left(\frac{n_{11}}{n_{1.}} - \frac{n_{11}}{n} \right) = \left(\frac{n-n_{1.}}{n} \right)^2 \text{Var} \left(\frac{n_{11}}{n_{1.}} \right) = \frac{(1-p_{1.})^2 p_{.1} (1-p_{.1})}{n_{1.}}$$

$$\text{Var} \left(\frac{n_{.1} - n_{11}}{n} \right) = (1-p_{1.})^2 \text{Var} \left(\frac{n_{.1} - n_{11}}{n-n_{1.}} \right) = \frac{(1-p_{1.})^2 p_{.1} (1-p_{.1})}{n-n_{1.}},$$

woraus

$$\text{Var } w' = (1-p_{1.})^2 p_{.1} (1-p_{.1}) \left(\frac{1}{n_{1.}} + \frac{1}{n-n_{1.}} \right)$$

$$= \frac{1}{n} \frac{1}{p_{1.}} (1-p_{1.}) p_{.1} (1-p_{.1})$$

folgt. Schließlich standardisieren wir w' , indem wir es auf $z' = \frac{w'}{\sqrt{\text{Var } w'}}$ bringen, was nach einfachen Umformungen zu (28) und (15) führt:

$$\text{Var } w' = \frac{1}{n} \cdot \frac{n}{n_{1.}} \cdot \frac{n-n_{1.}}{n} \cdot \frac{n_{.1}}{n} \cdot \frac{n-n_{.1}}{n}$$

$$= \frac{1}{n_{1.}} \cdot \frac{1}{n^3} (n \cdot n_{1.}) \cdot n_{.1} \cdot (n - n_{.1})$$

$$z' = \frac{w'}{\sqrt{\text{Var } w'}} = \frac{n \cdot n_{11} - n_{1.} \cdot n_{.1}}{n \cdot n_{1.} \cdot \sqrt{\frac{1}{n_{1.}} \cdot \frac{1}{n^3} (n - n_{1.}) \cdot n_{.1} \cdot (n - n_{.1})}} =$$

$$= \sqrt{n} \frac{n n_{11} - n_{1.} \cdot n_{.1}}{\sqrt{n_{1.} \cdot n_{.1} (n - n_{1.}) (n - n_{.1})}}$$

wobei berücksichtigt ist, daß $E(w') = p_{1.} \cdot p_{.1} = 0$ unter H_0 gilt.

Typ B

In Typ A (Homogenitätstest) haben wir feste Zeilensummen vorausgesetzt. Beim Unabhängigkeitstest arbeiten wir nur mit einer Stichprobe und setzen alle Randsummen als fest voraus. Faktisch sind die Randsummen zwar Zufallsvariablen, aber bei der Ableitung des Chi-Quadrat-Tests werden sie geschätzt - was zur Verringerung der Freiheitsgrade führt - und dann die n_{ij} als bedingte Zufallsvariablen in Abhängigkeit von den Randsummen untersucht. Auch die Untersuchung der Residuale beginnt so.

Die Werte $n_{i.}$ und $n_{.j}$ werden also bei der Ableitung als fest angenommen, und

$$p_{ij} = p_{i.} \cdot p_{.j} = \frac{n_{i.} \cdot n_{.j}}{n^2}$$

Das Kriterium wird als der Unterschied der beobachteten und der erwarteten Häufigkeit definiert, d.h.

$$w^* = \frac{n_{ij}}{n} - \frac{n_{i.} \cdot n_{.j}}{n^2} \quad (30)$$

Wir leiten das Resultat wiederum für das Feld (1,1) ab. Die Häu-

figkeiten f_{11} und w^* sind bei hinreichend großen n und n_{11} normal verteilt. Unter der Nullhypothese ist

$$E(f_{11}) = p_{11}$$

und $E(w^*) = 0$.

Um w^* standardisieren zu können, müssen wir seine Varianz ableiten. Wir untersuchen das Verhalten von f_{11} unter der Annahme, daß $n_{.1}$ und $n_{1.}$ feste Werte sind. Bezeichnen wir

$$X = n_{11}$$

$$Y = n_{1.} - n_{11}$$

$$Z = n_{.1} - n_{11}$$

und suchen wir die Verteilung von X unter der Bedingung, daß sowohl $X+Z = V$ als auch $X+Y = U$ feste Werte sind. Wenn X, U, V eine dreidimensionale Normalverteilung haben (was in unserem Fall asymptotisch gilt), dann hat das X unter der Bedingung $U = E(U)$ und $V = E(V)$ eine eindimensionale Normalverteilung mit der Erwartung

$$E(X|U, V) = E(X) = np_{11} = np_{1.} \cdot p_{.1} \quad (31)$$

und Varianz

$$\text{Var}(X|U, V) = \text{Var } X - [\text{Cov}(X, U), \text{Cov}(X, V)] \cdot \begin{bmatrix} \text{Var } U, \text{Cov}(U, V) \\ \text{Cov}(U, V), \text{Var } V \end{bmatrix}^{-1} \cdot \begin{bmatrix} \text{Cov}(X, U) \\ \text{Cov}(X, V) \end{bmatrix} \quad (32)$$

In (32) besteht der zweite Ausdruck rechts aus Matrixmultiplika-

tion und Matrixinversion. Er kann auch folgendermaßen dargestellt werden:

$$\frac{1}{\text{Var } U \cdot \text{Var } V - \text{Cov}^2(U, V)} \cdot [\text{Cov}(X, U), \text{Cov}(X, V)] \begin{bmatrix} \text{Var } V, -\text{Cov}(U, V) \\ -\text{Cov}(U, V), \text{Var } U \end{bmatrix} \cdot \begin{bmatrix} \text{Cov}(X, U) \\ \text{Cov}(X, V) \end{bmatrix} \quad (33)$$

Die einzelnen Größen folgen aus den Eigenschaften der unter H_0 angenommenen Multinomialverteilung:

$$\begin{aligned} \text{Var } X &= np_{11}(1-p_{11}) = np_{1.}p_{.1}(1-p_{1.}p_{.1}) \\ \text{Var } U &= np_{1.}(1-p_{1.}) \\ \text{Var } V &= np_{.1}(1-p_{.1}) \\ \text{Cov}(X, V) &= \text{Cov}(X, X+Z) = \text{Var } X + \text{Cov}(X, Z) = \\ &= np_{11}(1-p_{11}) - n \cdot p_{11}(p_{.1}-p_{11}) \\ &= np_{11}(1-p_{.1}) = np_{1.}p_{.1}(1-p_{.1}), \end{aligned}$$

da die Kovarianz zweier Zufallsgrößen X und Z aus einer Multinomialverteilung gleich

$-n \cdot p_x \cdot p_z$ ist.

$$\begin{aligned} \text{Cov}(X, U) &= \text{Cov}(X, X+Y) = \text{Var } X + \text{Cov}(X, Y) = \\ &= np_{11}(1-p_{1.}) = \\ &= np_{1.}p_{.1}(1-p_{1.}) \\ \text{Cov}(U, V) &= \text{Cov}(X+Y, X+Z) = \\ &= \text{Var } X + \text{Cov}(X, Y) + \text{Cov}(X, Z) + \text{Cov}(Y, Z) = \\ &= n(p_{11}-p_{1.}p_{.1}) = 0. \end{aligned}$$

setzt man diese Ausdrücke in (32) ein, so erhält man nach Umordnung

$$\text{Var } w = \text{Var}(X|U, V) = np_{1.}p_{.1}(1-p_{1.})(1-p_{.1}) \quad (34)$$

Schätzt man $\hat{p}_{1.} = \frac{n_{1.}}{n}$ und $\hat{p}_{.1} = \frac{n_{.1}}{n}$ aus der Stichprobe und setzt in die standardisierte Form

$$z^* = \frac{w^* - E(w^*)}{\sqrt{\text{Var}(w^*)}} \quad \text{ein,} \quad (35)$$

so erhält man leicht (28). Die Testgröße ist also gleich, ohne Rücksicht darauf, aus welchem Modell sie entstand. Die beiden Statistiken haben nur asymptotisch gleiche Eigenschaften. Im Fall A gibt es (C-1) unbekannte Parameter, die man mit Hilfe empirischer Daten schätzt, im Fall B gibt es (R+C-2) Parameter. Daher sind die Schätzungen und die Tests beim Typ A bei gleicher Tafelgröße und bei gleicher Anzahl der Beobachtungen exakter als beim Typ B.

7. SCHLUß

In dieser Arbeit wollten wir das korrekte Verfahren zur Aufstellung des Zeichenschemas für die Struktur der Beziehungen oder Differenzen in der Kontingenztafel einführen und beweisen. Die Motivation war zweifach: (a) Das Interesse für das Zeichenschema, das sich in der Praxis als sehr nützlich erwiesen hat, anzuregen, (b) das richtige Verfahren zur Konstruktion des Schemas zu zeigen. Es wäre schade, wenn die Idee der Konstruktion des Zeichenschemas vergessen würde, aber andererseits ist es nötig, daß das Programmieren auf theoretisch richtigen Grundlagen beruht. Wir gingen von der Arbeit Linhart/Šafár (1967) aus. Man kann jedoch das Zeichenschema aufgrund ganz anderer statistischer

Prinzipien bestimmen. Zum Schluß möchten wir betonen, daß diese praktische Methode ihre Beschränkung hat und die Resultate des Schemas auf keinen Fall absolutisiert werden können.¹⁴⁾

Anmerkungen

1. Dies ist die übliche Bezeichnung in den Kontingenztafeln. Es gilt

$$n_{i.} = \sum_{j=1}^C n_{ij}; n_{.j} = \sum_{i=1}^R n_{ij}; n = \sum_{i=1}^R \sum_{j=1}^C n_{ij}.$$

2. Diese Voraussetzung wird in der Soziologie meistens nicht erfüllt. Deshalb kann man die wahrscheinlichkeitstheoretischen Implikationen der Methode nicht ganz wörtlich nehmen; sie dienen lediglich als Leitfaden zur Interpretation. Man kann die Methode auf die Fragestellungen A und B natürlich auch unter der Voraussetzung anwenden, daß die Grundgesamtheiten als Ganze untersucht werden. Dabei nimmt man an, daß die Gesamtheiten durch einen Zufallsprozeß entstanden sind (z.B. Auswahl von Personen für Werkstätten).
3. Zur Illustration zeigen wir die Bezeichnungen z.B. der absoluten Häufigkeiten: $n_{11} = 236$; $n_{1.} = 1133$; $n_{.1} = 395$; $n = 1995$; der Prozente: $100f_{2/1}\% = 32.7\%$; $100f_{4/2}\% = 17.5\%$; $100f_{.3}\% = 30.9\%$; $100f_{1.}\% = 56.8\%$; $100f_{2.}\% = 43.2\%$.
4. Auch hier gilt die Forderung einer einfachen Zufallsstichprobe und die Anmerkung 2.
5. Wie bei der Tabelle 1 illustrieren wir die Bezeichnungen: $n_{23} = 103$; $n_{2.} = 210$; $n_{.3} = 413$; $n = 1097$; $100f_{23}\% = 9.4\%$; $100f_{2.}\% = 19.1\%$; $100f_{.3}\% = 37.6\%$.
6. Der Test ist begründet auf dem standardisierten Vergleich der beobachteten und der erwarteten Häufigkeiten in einzelnen Feldern nach dem Pearsonschen Prinzip

$$\chi^2 = \sum \frac{(\text{beobachtete Häufigkeit} - \text{erwartete Häufigkeit})^2}{\text{erwartete Häufigkeit}}$$

Die berechnete Statistik bezeichnet man als χ^2 und vergleicht sie mit der theoretischen Chi-Quadrat-Verteilung für die gegebene Anzahl der Freiheitsgrade.

- Wir bringen verschiedene Formeln, die aber mathematisch äquivalent sind. Die Wahl hängt von mehreren Faktoren ab: (a) Vom Datentyp, den wir vom Rechner oder aus Publikationen zur Verfügung haben (absolute oder relative Häufigkeiten oder diverse Kombinationen) (b) vom Programmierungstyp, bzw. vom Taschenrechner und seiner Logik; (c) von der persönlichen Vorliebe des Benutzers. Die angegebenen Formeln stellen nicht alle Möglichkeiten dar. Manchmal ist es vorteilhaft, die Formeln so zu gestalten, daß man bestimmte notwendige Zwischenergebnisse zur Verfügung hat usw. Bei den Berechnungen verfährt man am besten nach folgenden Prinzipien: (a) Ohne Rücksicht auf die Vereinfachung und Rundung der Resultate bei der Veröffentlichung sollen die Zwischenergebnisse die höchstmögliche numerische Exaktheit haben; (b) man soll das Verfahren wählen, bei dem man möglichst wenige Rundungen in Zwischenergebnissen durchführt (vor allem beim Dividieren); (c) Die Rechnung soll kontinuierlich durchgeführt werden, wobei die Zwischenergebnisse in Speichern aufbewahrt werden sollen; (d) Alle Rechnungen sollen grundsätzlich zweimal durchgeführt werden, wenn möglich von zwei Personen und unabhängig voneinander.
8. Die Ansichten über die Anwendbarkeit dieses Tests sind recht unterschiedlich. Im Vergleich mit der grundlegenden Arbeit (Cochran 1954) wurden die Forderungen ziemlich gelockert. Eine endgültige Entscheidung wurde noch nicht getroffen. Eine abschließende Auswertung wird offensichtlich auf weiteren umfangreichen Simulationen der Kontingenztafeln mit der Monte-Carlo-Methode und auf der Untersuchung der gewonnenen empirischen Verteilungen der Testcharakteristika beruhen. Ebenso unterscheiden sich auch die Ansichten über die Modifizierung

des Tests bei kleinen Stichproben mit Hilfe der sogenannten Korrektur für Kontinuität in 2x2-Tafeln. Pirie/Hamdan (1972) bringen eine Verbesserung der Korrektur für Kontinuität für beide Typen (A und B) der Kontingenztafeln. Grizzle (1967) empfiehlt, die Korrektur für Kontinuität nicht zu benutzen, da sie zur Konservativität der Tests beiträgt, besonders bei kleinen Stichproben. Bei kleinen Häufigkeiten in 2x2 Tafeln benutzt man den exakten Test von Fisher, für den es ausführliche Tabellen gibt (keine Berechnungen sind nötig, das Resultat findet man in den Tabellen direkt aufgrund der Verteilung der Häufigkeiten.) Für große Tafeln ($FG > 30$) und kleine erwartete Häufigkeiten siehe Cochran (1954). In Anbetracht dieser Uneinheitlichkeit und des Fehlens einer endgültigen Lösung empfehlen wir, die oben genannten Beschränkungen einzuhalten. Wir sind aber überzeugt, daß man die erste Forderung so lockern kann, daß sie der zweiten Forderung für 2xC Tafeln analog ist. Im Zweifelsfall empfiehlt es sich, den Rat eines Statistikers einzuholen. Beim Lösen wiederholter Aufgaben in Tafeln spezieller Typen empfehlen wir, die Monte-Carlo-Methode anzuwenden, um festzustellen, ob der Test anwendbar bzw. zu modifizieren ist.

9. Bei den meisten Rechnern ist dieses Verfahren am geeignetsten, da es einfach ist, Produkte zu summieren. Auch numerisch ist dieses Verfahren vorteilhaft.
10. Das eigentliche Ziel dieses Aufsatzes ist es, das korrekte Verfahren für das Zeichenschema zu erläutern. Moderne Rechen-systeme benutzen es nicht. Da sich das Verfahren bewährt hat, wäre es schade, wenn diese gute Idee tschechoslowakischer Methodologen in Vergessenheit geraten würde. Die Berechnung des z-Wertes ist nicht nur für die Konstruktion des ganzen Zeichenschemas von Bedeutung, sondern vor allem als Test für die Abweichung in einem Feld. Dann erhält das Testresultat seine ursprüngliche und richtige probabilistische Interpretation. Je nach Problem wählt man den einseitigen oder den zweiseitigen Test. Für den zweiseitigen Test sind die kritischen Werte von $|z_{ij}| < z_{\alpha}$ in Fußnote 11 angegeben. Für den einseitigen Test (d.h. die im voraus spezifizierte Richtung

des Unterschiedes) sind die kritischen Werte

α	z_{α}
0.1	1.2816
0.05	1.6449
0.01	2.3263
0.005	3.0902

11. Beim zweiseitigen Test sind die kritischen Werte

α	$z_{\frac{\alpha}{2}}$
0.05	1.9600
0.01	2.5758
0.001	3.2905

In der Praxis könnte man die Grenze auf 2.0, 2.5 und 3.0 vereinfachen, wodurch kein grober Fehler begangen würde. Diese Vereinfachung kompliziert jedoch die Situation dort, so man im Programm das Schema mit Hilfe der exakten Binomialverteilung ermitteln will (d.h. im Fall, daß $np_i.p_j < 5$). Will man $\alpha = 0.1$ benutzen, so ist $z = 1.6449$.

12. Für den Druck des Zeichenschemas kann man andere graphische Darstellungen wählen. Man kann die Zahl der Zeichen auf (++,+ ,0,-,--) beschränken oder auf vier (oder mehr) in beiden Richtungen erweitern: (+++,++ ,++ ,+ ,0,-,--,---,----). Wir vermuten jedoch, daß der Vorschlag von Linhart/Šafár (1967) für die analytische Arbeit optimal ist. Wir empfehlen auch die Praxis einiger Programme, die z-Werte auszudrucken.
13. Formel (15) ist etwas modifiziert für die Berechnung. Bei Rechnern mit Speicher ist es vorteilhaft, die Wurzel im Nenner für z_{11} zu speichern und bei der Berechnung der z-score alle Felder der ersten Zeile zu benutzen. Wenn die Zahl der Zeilen größer als die der Spalten ist, so speichert man den analogen Ausdruck für Spalten.
14. Während der Artikel in Druck war, haben die Autoren festgestellt, daß die Formeln für Residualabweichungen (ohne Ab-

leitung) in der Arbeit von Haberman (1973), und zwar im Kontext einer graphischen Methode für die Analyse der Residuen in Kontingenztafeln, enthalten sind.

Literatur

- Cochran, W.G., Some Methods for Strengthening the Common χ^2 Tests. Biometrics 10, 1954, 417-451.
- Grizzle, J.E., Continuity Correction in the χ^2 -test for 2x2 Tables. The American Statistician 21 (4), 1967, 28-32.
- Haberman, S.J., The Analysis of Residuals in Cross-Classified Tables. Biometrics 29, 1973, 205-220.
- Janko, J., Statistické tabulky. Praha, Nakladatelství ČSAV 1958.
- Kelley, T.L., Statističeskije tablicy. Moskva, Vyčislitelnyj centr AN SSSR 1966.
- Lancaster, H.O., The chi-squared Distribution. New York, Wiley 1969.
- Lewontin, R.C., Felsenstein, J., The Robustness of Homogeneity Tests in 2xN Tables. Biometrics 21, 1965, 19-33.
- Linhart, J., Šafár, Z., Programování třídění a statistických výpočtů pomocí samočinných počítačů. Sociologický časopis 4, 1967, 435-443.
- Pirie, W.R., Hamdan, M.A., Some Revised Continuity Corrections for Discrete Distributions. Biometrics 28, 1972, 693-701.
- Rao, C.R., Linear Statistical Inference and its Applications. New York, Wiley 1973 (2nd edition)
- Řehák, J., Řeháková, B., Měření statistické závislosti nominálních znaků. Sociologický časopis 4, 1973, 404-418.

DISKUSSION ZU ŘEHÁK & ŘEHÁKOVÁ, ANALYSE VON KONTINGENZTAFELN

W. Matthäus

Wenn man das Vorzeichenmuster als Ganzes beschreiben will und somit eine zusammengesetzte Aussage über die Gesamtheit der Zellen macht, wie auf Seite 12 bei Beisp 1 2 geschehen, - möglicherweise um dieses Muster in einer theoretischen Argumentation zu verwenden - so befindet man sich in einer ähnlichen Lage, wie wenn man bei der Varianzanalyse nach einem signifikanten Haupteffekt multiple Vergleiche durchführt, um den globalen Effekt zu spezifizieren. Wenn man das Fehlerrisiko der zusammengesetzten Aussage über das Unterschieds- oder Abhängigkeitsmuster nicht übermäßig anschwellen lassen will, so muß man das Signifikanzniveau bei den einzelnen Teilaussagen adjustieren, d.h. verschärfen. Eine obere Grenze des Fehlerrisikos der zusammengesetzten Gesamtaussage kann man mit Hilfe der Bonferroni-Ungleichung abschätzen (R.G. MILLER 1966: 6-8, 67-70). Für die Wahrscheinlichkeit einer fälschlichen Ablehnung der zusammengesetzten H_0 über k Zellen ("in keiner der k Zellen gibt es einen signifikanten Unterschied zwischen Erwartung und Beobachtung") gilt: $P(F) \leq \alpha_1 + \dots + \alpha_k$. Sollen alle k Signifikanzniveaus gleich festgesetzt werden, sagen wir auf α' , und soll das Gesamtrisiko nicht größer als α sein, so muß man $\alpha' = \alpha/k$ setzen. Dann wird $P(F) \leq k (\alpha/k) = \alpha$. Für jede der k Zellen hat man also das Signifikanzniveau auf α/k statt auf α zu setzen. Wenn z.B. ein Zeichenschema für 6 Zellen erstellt werden soll, testet man jede Zelle nicht auf dem Niveau 0.05, sondern auf dem Niveau $0.05/6 \approx 0.01$.

Die Frage der Interpretationseinheit - Einzelaussage oder zusammengesetzte Aussage über ein Muster - ist allerdings umstritten. Es hat sich noch keine Konvention als Entscheidungshilfe durchgesetzt.

G. Altmann

Der Test für eine Zelle (§6) läßt sich auch auf eine andere Weise ableiten. Kendall und Stuart (1960, II: 551) haben gezeigt, daß die

Verteilung einer Zelle der 2 x 2 Tafel mit festen Randverteilungen hypergeometrisch ist. Es läßt sich zeigen, daß dies auch für eine RxC-Tafel gilt.

Bezeichnen wir $f(n_{11}, n_{12}, \dots, n_{rc}) = f(n_{ij})$ als die gemeinsame Verteilung aller Zellenvariablen und $f(n_{1.}, n_{2.}, \dots, n_{r.}, n_{.1}, n_{.2}, \dots, n_{.c}) = f(n_{i.}, n_{.j})$ als die gemeinsame Verteilung aller Randvariablen. Wir suchen die bedingte Verteilung der Zellenvariablen, d.h. $f(n_{ij} | n_{i.}, n_{.j})$. Da die gemeinsame Verteilung $f(n_{ij}, n_{i.}, n_{.j}) = f(n_{ij})$ - weil bei bekannten n_{ij} ($i=1,2,\dots,r$; $j=1,2,\dots,c$) alle $n_{i.}$ und $n_{.j}$ gegeben sind - so ist

$$f(n_{ij} | n_{i.}, n_{.j}) = \frac{f(n_{ij})}{f(n_{i.}, n_{.j})}.$$

Unter Nullhypothese, daß $p_{ij} = p_{i.} p_{.j}$ ist, folgt

$$f(n_{ij}) = \frac{n!}{\prod_{i,j} n_{ij}!} \prod_{i,j} p_{ij}^{n_{ij}} = \frac{n!}{\prod_{i,j} n_{ij}!} \prod_i p_{i.}^{n_{i.}} \prod_j p_{.j}^{n_{.j}}$$

$$\text{und } g(n_{i.}, n_{.j}) = \frac{n!}{\prod_i n_{i.}!} \prod_i p_{i.}^{n_{i.}} \frac{n!}{\prod_j n_{.j}!} \prod_j p_{.j}^{n_{.j}},$$

woraus sich

$$f(n_{ij} | n_{i.}, n_{.j}) = \frac{\prod_i n_{i.}! \prod_j n_{.j}!}{n! \prod_{i,j} n_{ij}!} \quad (1)$$

ergibt.

Für eine gegebene Zelle (i,j) ergibt sich aus (1) leicht

$$E(n_{ij}) = \frac{n_{i.} n_{.j}}{n} \quad (2)$$

$$\text{und } V(n_{ij}) = \frac{n_{i.} n_{.j} (n - n_{i.}) (n - n_{.j})}{n^2 (n-1)}, \quad (3)$$

$$\text{so daß } z = \frac{n_{ij} - \frac{n_{i.} n_{.j}}{n}}{\sqrt{\frac{n_{i.} n_{.j} (n - n_{i.}) (n - n_{.j})}{n^2 (n-1)}}} \quad (4)$$

asymptotisch normalverteilt ist, $N(0,1)$. Formel (4) ist mit (14) und (28) asymptotisch identisch. Der einzige Unterschied besteht in der Größe (n-1) statt n bei Řehák/Řeháková. Bezeichnet man aber $\frac{n_{.j}}{n} = p$, so sieht man, daß (2) und (3) eben die Momente der hypergeometrischen Verteilung darstellen, nämlich

$$E(n_{ij}) = n_{i.} p \quad (2a)$$

$$V(n_{ij}) = n_{i.} p q \frac{(n - n_{i.})}{n-1}. \quad (2b)$$

Sind die einzelnen Konvergenzbedingungen erfüllt, so kann man also die Tests für die Abweichung in einer gegebenen Zelle auch mit Hilfe der Binomialverteilung bzw. Poissonverteilung bzw. Normalverteilung durchführen.

J. Řehák, B. Řeháková

Both notes contributed by the editors are of great importance. While the latter gives a different mathematical background to the method and a guideline for practical computation in small samples, the former formulates the important methodological problem of simultaneous inference.

Matthäus shows how it is possible to adjust the significance levels to get a probabilistically correct structure of displayed signs. In research practice we always have to ask a fundamental

question while drawing several parallel statistical inferences: "Are we interested in each result separately or do we accept the whole structure of all inferences as a final statement?" In the latter case the probability of a conclusion is much more complicated. For a given case of cells in a contingency table and their departures from independence, we might use the usual significance levels in a special situation when we are interested in departures in a priori determined cells where each cell represents an independent problem. If we want to accept an overall structure of departures, we can use the suggestion of Matthäus, but one should take the following into account: he suggests adjusting the significance level by taking the number of cells for k (the adjustment factor). We would propose replacing the number of cells by the number of degrees of freedom. This is motivated as follows: in a 2×2 table the chi-square statistic is equivalent to the testing of any cell, and vice versa: any cell departure gives a significant chi-square statistic. The square of the z -score for any cell is equal to the chi-square statistic. Moreover, any sign determines all the other three. By using $k = 4$, i.e. the number of cells, we would have an obviously incorrect inference. In this case indeed $k = 1$, i.e. the number of degrees of freedom. In general we can take k as a number of linearly independent statements. The rest of the statements can be derived from these basic ones. In the case that we are interested only in an a priori given subset of cells, the k should be determined especially from the cells' configurations. In some cases it is quite a tricky problem.

We have another comment to the use of a sign scheme based on the Bonferroni inequality. The sign scheme method has a primarily exploratory role in our understanding. The more stringent simultaneous significance levels will give a reduced structure of signs which can be probabilistically accepted but at the cost of losing information we could use for generating hypotheses. The unadjusted significance levels on the other hand give richer and more suggestive structure but involve the danger that random fluctuations may be interpreted as substantive. Our experience however shows that a researcher with good knowledge of a problem can orient himself in the structure quite well and filter out

the randomly caused signs by their meaning. The process of interpretation resembles the interpretation of factor analytic results and possesses all aspects of exploratory data analysis with its dangers and advantages. For practical work and programming we propose to print the scheme of signs as given in the paper (on the basis of the ordinary significance levels) and to give the critical value for z -scores on the basis of adjustment by df . This would give the possibility of accepting the probabilistically relevant structure and keeping the richer graphic display.

Altmann's note gives a different important insight into the problem that has been solved by routine methods of a standard theory in our paper. This suggestion leads to an exact Fisher test for a cell against the rest of the table. It introduces the possibility of an exact solution for a case of small n . We believe that this is going to be relevant in large tables where observations are scattered so that frequencies are small (we are frequently interested only in a part of a table). Asymptotically there will be no difference between the two formulas. For small n , the exact Fisher test should be used routinely as a formula for significance computations. Recently one of us in fact (B.Ř.) used this way of preparing an analytical package. Altmann's approach also sheds some light on results that might seem strange to some readers, namely to the fact that the two different situations of type A and B are solved by the same formulas and that the derivation of a case where the marginals are obtained at random is based on the contradictory assumption that they are conditioned and therefore held fixed. We refer to the discussion of Fisher's test, e.g. in Kendall, Stuart (1960). The case of a larger table is analogous. The equivalence of the two procedures is asymptotic (i.e. for large n). Fixing the marginals in type B is necessary in order to be able to estimate marginal probabilities to obtain the expected values in individual cells.

Literatur

- Kendall, M.G., Stuart, A., The advanced theory of statistics. London, Griffin 1960. Vol. II.
Miller, R.G., Simultaneous statistical inference. New York, McGraw-Hill 1966.

EINIGE KRITISCHE BEMERKUNGEN ZU W. LEHFELDT "ZUR NUMERISCHEN ERFASSUNG DER SCHWIERIGKEIT DES SPRECHBEWEGUNGSABLAUFS"

R. Grotjahn, Bochum

Wie das Beispiel anderer Wissenschaften zeigt, wird in vielen Fällen erst durch die Weiterentwicklung geeigneter Meßmodelle die Lösung bestimmter Probleme ermöglicht. Es ist deswegen begrüßenswert, daß Werner Lehfelddt (1980) in seinem in Glottometrika 2 erschienenen Artikel "Zur numerischen Erfassung der Schwierigkeit des Sprechbewegungsablaufs" den Versuch unternommen hat, einige bei der Quantifizierung der Schwierigkeit des Sprechbewegungsablaufs auftretende methodologische Probleme näher zu analysieren. Um die methodologische Diskussion über die meßtheoretischen Grundlagen der Analyse der Schwierigkeit des Sprechbewegungsablaufs weiterzuführen, sollen nun einige m.E. problematische Punkte in den Ausführungen Lehfelddts kurz erwähnt werden. Der an einer Begründung der geäußerten Kritik sowie an einer weiterführenden Diskussion des Problems der Quantifizierung der Schwierigkeit des Sprechbewegungsablaufs interessierte Leser sei auf einen in Kürze erscheinenden Artikel verwiesen (vgl. Grotjahn 1980).

Die Kritik an Lehfelddt läßt sich in die Problemkreise 'begriffliche Klarheit' sowie 'Indexbildung und Validierung' unterteilen.

(a) Begriffliche Klarheit. Die Verwendungsweise des Begriffs der Schwierigkeit, der in den Ausführungen Lehfelddts eine zentrale Stellung einnimmt, ist relativ unpräzise. Zum einen wird Schwierigkeit als Synonym zu Kompliziertheit verwendet, ohne daß Kompliziertheit ihrerseits definiert wird (vgl. z.B. S. 47 u. 58). Zum anderen wird die Schwierigkeit von Lautverbindungen anscheinend auch weitgehend mit der Verbindlichkeit und dem Präzisionsgrad von Bewegungstypen gleichgesetzt, oder es wird zumindest ohne empirische Überprüfung unmittelbar von den Variablen Verbindlichkeit und Präzisionsgrad auf die Variable Schwierigkeit geschlossen (vgl. z.B. S. 52f). Eine Problematisierung der Beziehung zwischen dem subjektbezogenen Begriff der Schwierigkeit und den von Lindner (1975) phonetisch-systemimmanent defi-

- 35 -

nierten Begriffen der Verbindlichkeit sowie des Präzisionsgrades erfolgt nicht.

(b) Indexbildung und Validierung. Die theoretische Begründung der Gewichtung der einzelnen Bewegungstypen erscheint wenig befriedigend. Die Indexbildung erfolgt auf der Basis eines additiven Modells, das vom Standpunkt der Repräsentationsmessung Intervallskalenniveau voraussetzt. Die verwendeten Gewichte haben jedoch bestenfalls Ordinalskalengualität. Der Index berücksichtigt nicht eine mögliche Korrelation zwischen Präzisionsgrad und Organtyp. Es werden nur einige wenige Indexwerte berechnet. Wie jedoch eine Überprüfung der Verteilung des Indexes anhand einer größeren Anzahl von Lautverbindungen zeigt, weist der Index keine zufriedenstellende Trennschärfe auf. Eine Validierung und damit auch ein Nachweis der theoretischen Fruchtbarkeit oder der praktischen Nützlichkeit des vorgeschlagenen Indexes findet nicht statt. Validiert man den Index anhand des Außenkriteriums 'Anzahl der Stammelfehler deutscher Grundschüler bei zweigliedrigen Lautverbindungen', ergibt sich eine Nullkorrelation. Es bleibt somit offen, was der Index überhaupt mißt. Akzeptiert man die Validität des Außenkriteriums 'Anzahl der Stammelfehler', dann dürfte der Index jedenfalls kaum die Schwierigkeit des Sprechbewegungsablaufs messen.

Literatur

- Grotjahn, R., Zur Quantifizierung der Schwierigkeit des Sprechbewegungsablaufs. In: Grotjahn, R./Hopkins, E. (Eds.), Empirical Research on Language Teaching and Language Acquisition. Bochum, Brockmeyer 1980
- Lehfelddt, W., Zur numerischen Erfassung der Schwierigkeit des Sprechbewegungsablaufs. In: Grotjahn, R. (Ed.), Glottometrika 2. Bochum, Brockmeyer 1980, 44-61
- Lindner, G., Der Sprechbewegungsablauf. Eine phonetische Studie des Deutschen. Berlin, Akademie-Verlag 1975.

Die vorliegende Arbeit¹ untersucht anhand von musikalischem Textmaterial (d.h. Musikwerken als Textgebilden) aus verschiedenen Stilen die Prinzipien der Rekurrenz kleiner motivartiger melodischer Elemente im Text. Diese Organisationsprinzipien, die für Texte verschiedenster Stile *allgemeingültig* sind, werden mit Hilfe von Methoden beschrieben, die in der quantitativen Linguistik entwickelt worden sind. Es wird gezeigt, daß zwischen diesen Prinzipien und dem musikalischen Text *als Ganzem* ein Zusammenhang besteht und daß hier analoge Verhältnisse wie bei den Organisationsprinzipien der Rekurrenz von Wörtern in literarischen Texten vorliegen.

1.

Die Untersuchungen musikalischer Texte aus verschiedenen Stilen zeigen, daß die Wiederholung einen der Hauptfaktoren der musikalischen Formenbildung und überhaupt das wichtigste Mittel zur Gestaltung des "musikalischen Inhalts" darstellt. Das Prinzip der Wiederholung durchzieht den musikalischen Text auf sämtlichen formalen Ebenen - von ganzen Sätzen musikalischer Werke bis hin zu Intervallen und rhythmischen Einheiten.

Aus dieser Sicht kann die musikalische Form auf jeder beliebigen Ebene betrachtet werden als eine organisierte Aufeinanderfolge von "neuen", d.h. vom Textanfang an erstmalig vorkommenden Elementen und solchen Elementen, die bereits aufgetretene wiederholen. Die Erforschung der Gesetzmäßigkeiten, nach denen sie organisiert ist, stellt für die Musiktheorie eine wichtige Aufgabe dar. Wie die bisherigen Ergebnisse zeigen, ist es allerdings mit herkömmlichen musikwissenschaftlichen Methoden nur in unterschiedlichem Ausmaße möglich, die Organisation der im Text auftretenden Wiederholungen und Alternationen von *großen* Teilstücken (in der Größenordnung von Sätzen) einerseits und *kleinen* Teilstücken (vom Range eines Motivs) andererseits zu analysieren. Während für Großeinheiten viele Prinzipien ihrer Wiederholung und Alternation im Text wohlbekannt sind (teilweise sind sie in Formschemata fixiert), sind für *kleine* Elemente derartige Prinzipien noch nicht nachgewiesen.² Es ist unklar, ob die Rekurrenz kleiner

Elemente im musikalischen Text innerhalb des jeweiligen Stils (einer Epoche) einheitlich organisiert ist, ob sie bei jedem Komponisten individuell ausgeprägt ist, oder ob es letztlich ein für jeden Text gültiges durchgängiges Organisationsprinzip gibt. Eine Antwort hierauf wird möglich durch eine konkrete Untersuchung der Rekurrenzstruktur kleiner musikalischer Elemente in musikalischen Texten verschiedener Stile.

Für eine derartige Untersuchung ist offenbar folgendes notwendig:

- a) die Abgrenzung einer bestimmten Textebene (z.B. der melodischen),
- b) die Segmentierung des Textes auf dieser Ebene in Elemente eines bestimmten Typs;
- c) bestimmte Kriterien zur Unterscheidung der Elemente voneinander;
- d) die Auswertung des Textes auf der abgegrenzten Ebene bezüglich der Rekurrenz dieser Elemente.

Im Untersuchungsergebnis läßt sich dann jedem der ermittelten distinkten Textelemente ein bestimmter Zahlenwert zuordnen - seine Vorkommenshäufigkeit im Text. Dementsprechend kann dem ganzen Text eine Tabelle der Form

$$\begin{pmatrix} a_1, a_2, a_3, \dots, a_n \\ p_1, p_2, p_3, \dots, p_n \end{pmatrix} \quad (1)$$

zugeordnet werden, in der p_i die Vorkommenshäufigkeit des Elementes a_i im Text ist.

Die ermittelten Daten können unter zwei Gesichtspunkten analysiert werden. Einerseits lassen sich, bei einer quantitativen Interpretation der Elemente a_i , anhand einer Tabelle der genannten Form (1) statistische Charakteristiken berechnen, und ihre Werte können dann für unterschiedliche Texte und Textgruppen verglichen werden. Das ist das Verfahren der "statistischen Musikanalyse", das bei der Untersuchung der stilistischen Besonderheiten der Werke verschiedener Komponisten durchaus effektiv angewandt werden kann (vgl. DETLOVS 1968, FUCKS 1975, ROJTERŠTEJN 1973). Hingegen können metastilistische Prinzipien, nach denen die Rekurrenz von Elementen im musikalischen Text organisiert ist, bei einem solchen Verfahren schwerlich festgestellt werden, denn die Vorkommens-

häufigkeiten von Elementen, und somit die Werte der statistischen Charakteristiken, weichen in der Regel von Stil zu Stil erheblich voneinander ab. Darüber hinaus führt die Notwendigkeit einer zuverlässigen Bewertung dieser Häufigkeiten dazu, daß diese nicht jeweils für einzelne Texte berechnet werden, sondern für stilistisch homogene Mengen von Texten oder Textausschnitten, daß also die Geschlossenheit des Textes außer acht gelassen wird.

Es besteht andererseits jedoch auch die Möglichkeit, die Daten einer Tabelle (1) dahingehend auszuwerten, daß nicht Paare "Element-Häufigkeit" betrachtet werden, sondern die Häufigkeiten selbst und die Beziehungen zwischen ihnen. Dabei wird davon abstrahiert, welche Häufigkeit welchem Element entspricht. Ein derartiges Vorgehen erlaubt es, die Struktur der Elementwiederholungen im Text detailliert zu untersuchen: die Beziehungen zwischen "häufigen" und "seltenen" Elementen, zwischen rekurrenten und nichtrekurrenten Elementen etc. Hierbei entsteht nicht das Problem der zuverlässigen Bewertung der "Wahrscheinlichkeit" dieses oder jenes konkreten Elementes. Zudem können schließlich auch bei identischen Mengen von Vorkommenshäufigkeiten von Elementen diese Elemente selbst voneinander gänzlich verschieden sein. Daher ist es sehr wahrscheinlich, daß mittels der Analyse von Häufigkeitsmengen $\{p_i\}$ und der in ihnen geltenden Beziehungen *allgemeingültige Organisationsprinzipien der Elementrekurrenz im musikalischen Text als einem Ganzen nachgewiesen werden können*.

Zugunsten einer derartigen Hypothese sprechen, wenn auch indirekt, die Ergebnisse der Untersuchung eines ähnlich gelagerten Problems in der quantitativen Linguistik. Dort nämlich werden Mengen von Vorkommenshäufigkeiten von Elementen in lexikalischen Stichproben (man bezeichnet sie gewöhnlich als *statistische* oder *Häufigkeitsstruktur* einer Stichprobe) schon seit langem eingehend untersucht. Die Analyse der Häufigkeitsstruktur von Stichproben machte es möglich, Organisationsprinzipien der Wortrekurrenz nachzuweisen, die für verschiedenste Stichproben allgemeingültig sind (FRUMKINA 1964, ORLOV 1970).

So ließ sich zum Beispiel in jeder Stichprobe eine große Anzahl wenig frequenter Wörter feststellen und demgegenüber eine geringe Anzahl häufiger Wörter. Je kleiner dabei die Häufigkeit eines

Wortes in der Stichprobe war, desto größer war die Anzahl der Wörter mit dieser Häufigkeit. Insbesondere die einmal vorkommenden, d.h. in der Stichprobe nicht wiederholten Wörter bildeten annähernd die Hälfte ihres Lexikons. Wenn man nun weiterhin die Menge $\{p_i\}$ der relativen Vorkommenshäufigkeiten von Wörtern in der Stichprobe nach abnehmenden Werten ordnete, konnte der allgemeine Verlauf der Häufigkeitsabnahme durch den Ausdruck

$$p_i = \frac{K}{(B+i)^\gamma}, \quad i=1,2,3,\dots,n \quad (2)$$

näherungsweise berechnet werden. Hierbei ist p_i die Häufigkeit i -ter Ordnung in der geordneten Häufigkeitsmenge, n der Umfang des Lexikons der Stichprobe, und K , B und γ sind Konstanten. Den Ausdruck (2) bezeichnet man als Zipf-Mandelbrotsches Gesetz (vgl. FRUMKINA 1964, ORLOV 1970).

In jüngeren Arbeiten (ORLOV 1970, 1975) wurde gezeigt, daß derartige Gesetzmäßigkeiten an vollständigen Texten in sich abgeschlossener literarischer Werke am genauesten erfüllt werden. Dabei gilt $\gamma=1$, K und B sind Funktionen der Textlänge³ N_0 und der Frequenz des im Text am häufigsten vorkommenden Wortes $p_{\max}$, und der ganze Ausdruck für p_i nimmt damit folgende Form an⁴:

$$p_i = \frac{\frac{1}{\ln(N_0 p_{\max})}}{\frac{1}{p_{\max} \ln(N_0 p_{\max})} - 1 + i} \quad (3)$$

Der Umfang $n(N_0)$ des Lexikons des Textes, auf den das Gesetz (3) angewandt wird, und die Zahl $n_m(N_0)$ der einzelnen Wörter, von denen jedes im Text m -mal gebraucht ist, wird nach ORLOV (1975) durch die Beziehungen

$$n(N_0) = \frac{N_0 - \frac{1}{p_{\max}}}{\ln(N_0 p_{\max})} \quad (4)$$

$$n_m(N_0) = \frac{n(N_0)}{m(m+1)} \quad (5)$$

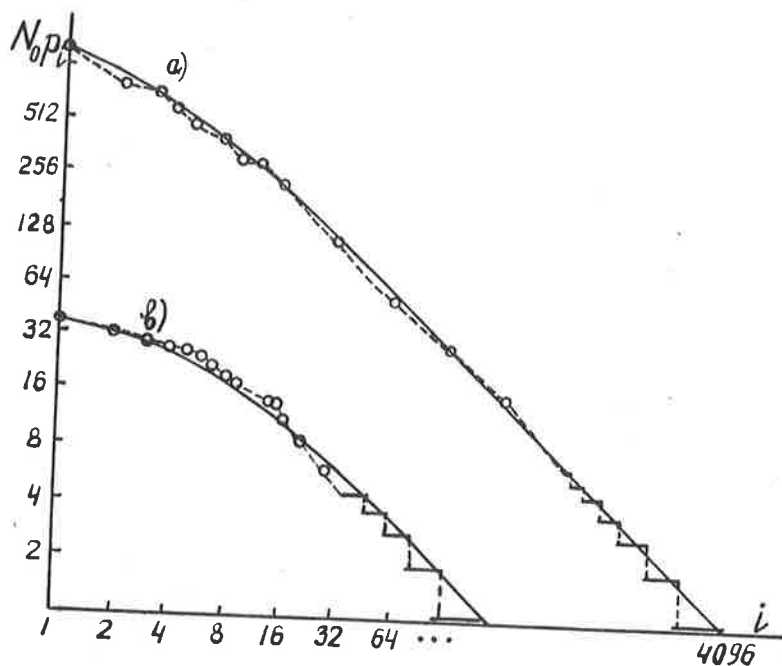


Abb. 1: Graphische Darstellung der Häufigkeitsstruktur literarischer Texte.

Experimentell ermittelte Kurve und theoretische Beschreibung nach dem verallgemeinerten Zipf-Mandelbrot'schen Gesetz (3) (entnommen aus Orlov 1975)

a) A.S. Puškin, Kapitanskaja dočka (Die Hauptmanns-tochter)

b) altruss. Byline Avdot'ja Rjazanočka
Die durchgezogene Linie verbindet die Punkte, die gemäß der theoretischen Beschreibung berechnet wurden, die gestrichelte Linie die experimentell ermittelten. Die "Stufen" im rechten Teil der graph. Darstellung sind bedingt durch die große Anzahl wenig frequenter, d.h. einmal, zweimal etc. vorkommender, Wörter im jeweiligen Text, der bezüglich der Geltung des Gesetzes (3) untersucht wurde.

definiert. Demnach ist die Struktur der Wiederholungen von Wörtern im literarischen Text abhängig von der Länge dieses Textes und von der Frequenz des in ihm am häufigsten vorkommenden Wortes. Je län-

ger ein Text ist, desto größer ist, bei annähernd gleichen Werten für p_{max} , der Umfang seines Lexikons. Einen bedeutenden Anteil am Lexikon eines Textes stellen die wenig frequenten Wörter, d.h. einmal, zweimal, dreimal etc. vorkommende. Die Zahl der m-mal vorkommenden Wörter nimmt ab mit dem Anwachsen von m.

Wie bei ORLOV (1970, 1975) gezeigt wird, beschreiben die Beziehungen (3), (4) und (5) die Struktur von Wortwiederholungen in vollständigen literarischen Texten verschiedener Stile (vgl. Abb.1). Für eine derartige Beschreibung ist die Vollständigkeit des untersuchten Textes von wesentlicher Bedeutung: für Textausschnitte und Textkorpora sind die Beziehungen (3), (4) und (5) in der Regel nicht erfüllt. So beträgt beispielsweise die Abweichung der tatsächlichen Werte für den Lexikon-Umfang aus (4) 100-150 %. Somit gehorcht die Rekurrenzorganisation von Wörtern im Text einer Reihe von allgemeinen Prinzipien und steht im Zusammenhang mit der vollen Textlänge. Die Analyse der Häufigkeitsstruktur hat es ermöglicht, diese Prinzipien nachzuweisen und sie mathematisch zu beschreiben.

Diese anhand von literarischen Texten ermittelten Ergebnisse legten den Gedanken nahe, mithilfe der Analyse von Häufigkeitsstrukturen auch die Organisation der Rekurrenz von Elementen in musikalischen Texten zu untersuchen, die zu diesem Zweck auf einer bestimmten Ebene jeweils in ihrer Gesamtheit, d.h. von Anfang bis Ende, betrachtet werden müßten.

2.

Wir waren also vor die Aufgabe gestellt, unter Berücksichtigung verschiedener Stile die Struktur von Elementwiederholungen im musikalischen Text zu untersuchen, und zwar auf der melodischen Ebene (in homophonen Texten anhand der "Oberstimme", in polyphonen anhand der Gesamtheit der Stimmen).

Bei einer derartigen Untersuchung tritt jedoch folgende Schwierigkeit auf. Wählt man als Elementareinheit eine Einheit mit fixierter Länge (ein melodisches Intervall, eine Aufeinanderfolge von Intervallen u.ä.), so führt das zu einer unnatürlichen Segmentierung des musikalischen Textes, die vergleichbar wäre mit der Zerstückelung eines literarischen Textes in Segmente von n Buch-

staben Länge. Einheiten mit variabler Länge hingegen, wie sie in der Musikwissenschaft schon bekannt sind (Motiv, Submotiv), sind nicht eindeutig definiert und werden uneinheitlich behandelt (vgl. KAC 1972, TJULIN 1974). Angesichts dieser Lage ergibt sich die Notwendigkeit, exakt definierte Elementareinheiten mit variabler Länge abzugrenzen, die auf bekannten musikalischen Gesetzmäßigkeiten beruhen.

Eine Einheit dieser Art ist das "formale Motiv" (F-Motiv), das wir schon in einer früheren Arbeit (BORODA 1973) definiert haben. Dabei sind wir von den in der Musikwissenschaft bekannten metrisch-rhythmischen Organisationsprinzipien taktgebundener Musik ausgegangen. Auf die Eigenschaften des F-Motivs und sein Verhältnis zu Motiv und Submotiv sind wir in einer weiteren Arbeit (BORODA 1977) auch bereits eingegangen. Im folgenden führen wir eine kurze Definition des F-Motivs an:

Als "formales Motiv" (F-Motiv) wird ein Segment der melodischen Linie bezeichnet, das innerhalb einer der nachstehenden vier elementaren metrisch-rhythmischen Gruppierungen einzugrenzen ist:

a) innerhalb eines vollständigen Minimaltaktes, d.h. von zwei bzw. im Triolenrhythmus drei gleichlangen Tönen, deren erster metrisch stärker ist als die übrigen (vgl. Abb. 2a);

b) innerhalb eines partiellen Minimaltaktes, d.h. eines Tones bzw. im Triolenrhythmus zweier gleichlanger Töne, die zusammen mit dem unmittelbar folgenden Ton noch keinen vollständigen Minimaltakt ergeben, da dieser Ton sich entweder von der Länge her unterscheidet oder metrisch stärker ist als die unmittelbar vorausgehenden (vgl. Abb. 2a);

c) innerhalb einer anwachsenden Tonfolge, d.h. einer Tonfolge, bei der jeweils der folgende Ton länger ist als der vorausgehende (vgl. Abb. 2b);

d) innerhalb einer metrisch-rhythmischen Minimalgruppierung, d.h. einer Verbindung aus einem (vollständigen oder partiellen) Minimaltakt und einer anwachsenden Tonfolge, die mit dem letzten Ton des Minimaltaktes beginnt (vgl. Abb. 2c).⁶

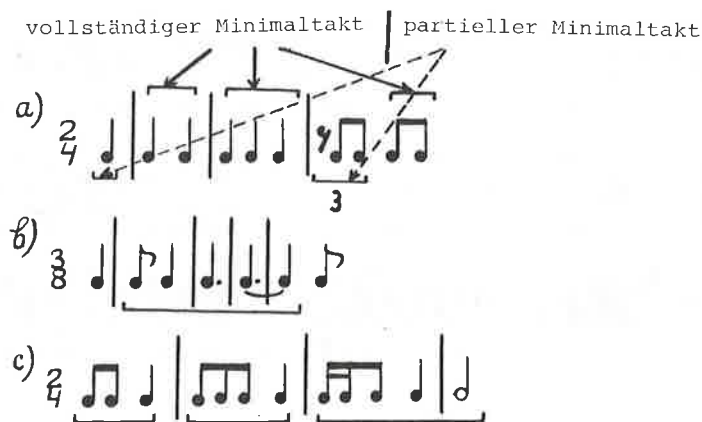


Abb. 2

Wie in den erwähnten Arbeiten (BORODA 1973, 1977) gezeigt wird, kann die melodische Linie eines musikalischen Textes mit Taktstruktur auf diese Weise von Anfang bis Ende eindeutig in F-Motive segmentiert werden. Einige Beispiele einer solchen Segmentierung zeigt Abb. 3.





Abb. 3: F-Motiv-Segmentierung von Melodieausschnitten musikalischer Texte aus verschiedenen Stilen.

(Die einzelnen F-Motive sind durch Klammern der Art  kenntlich gemacht.)

- a) J.S. Bach. Präludium Nr. 20 aus dem Wohltemperierten Klavier, 1. Teil BWV 865
- b) L. van Beethoven. Sonate für Klavier Nr. 10
- c) F. Schubert. Gretchen am Spinnrade D 118
- d) P.I. Čajkovskij. Symphonie Nr. 5 e-moll op. 64
- e) S.S. Prokof'ev. Romeo und Julia
- f) D.D. Šostakovič. Fuge für Klavier op. 87 Nr. 12

Aus diesen Beispielen und den genannten Definitionen geht hervor, daß das F-Motiv eine Elementareinheit mit variabler Länge ist. Die Segmentierung der melodischen Linie in F-Motive steht in engem Zusammenhang mit deren metrisch-rhythmischer Struktur. Im Ausnahmefall kann ein F-Motiv sich auch mit einem Motiv oder Submotiv decken (wie in den Beispielen b) und c) der Abb. 3).

Nachdem wir nun die Elementareinheit, das F-Motiv, festgelegt haben, können wir unsere Aufgabe exakter bestimmen: sie besteht also darin, für musikalische Texte aus verschiedenen Stilen die Organisation der Rekurrenz von F-Motiven im melodischen Text-"Schnitt" zu untersuchen, d.h. die Häufigkeitsstruktur und andere Charakteristiken dieser Organisation zu analysieren.

Als Untersuchungsmaterial wurden hierzu homophone Texte (mit einer signifikanten "Oberstimme") und polyphone Texte (mit untereinander melodisch gleichwertigen Stimmen) aus dem Stil-Spektrum vom 18. bis zum 20. Jahrhundert herangezogen. Der melodische "Schnitt" eines jeden Textes, d.h. die "Oberstimme" eines homophonen bzw. die Gesamtheit der Stimmen eines polyphonen Textes, wurde in F-Motive segmentiert,⁷ und es wurden für ihn folgende Charakteristiken bestimmt:

- 1) die Anzahl aller im melodischen "Schnitt" gebrauchten F-Motive (im folgenden bezeichnet als Textlänge);
- 2) die Anzahl aller nach einem bestimmten Kriterium distinkten F-Motive (im folgenden bezeichnet als Motivinventar des Textes⁸). Dabei galten zwei F-Motive ausschließlich für den Fall als identisch, wenn das eine aus dem anderen durch eine exakte sequenz-artige Transposition (also eine Parallelverschiebung in der Tonhöhe) unter Beibehaltung der jeweiligen Tonlängen gewonnen werden konnte;⁹
- 3) die Vorkommenshäufigkeiten eines jeden distinkten F-Motivs. Die so ermittelte Menge von Häufigkeiten wurde nach abnehmenden Werten geordnet und stellte somit die Häufigkeitsstruktur des Textes dar.

Die Analyse der Beziehungen zwischen diesen Charakteristiken für die verschiedenen Texte ließ folgende Gesetzmäßigkeiten erkennen:

- 1) In jedem Text ließ sich einerseits eine geringe Anzahl häufig wiederholter F-Motive feststellen und andererseits eine beträchtliche Anzahl seltener F-Motive, die in diesem Text nur einmal, zweimal etc. vorkamen. Die Anzahl solcher m-mal vorkommenden F-Motive wuchs mit abnehmendem m.
- 2) Es zeigte sich, daß das Motivinventar eines Textes mit der Länge dieses Textes in Zusammenhang steht. In Texten mit unterschiedlicher Länge war auch das Motivinventar unterschiedlich umfangreich (je länger der Text, desto umfangreicher).

Diese beiden eben genannten Gesetzmäßigkeiten gelten, wie wir feststellen konnten, für vollständige Texte. Bei Textausschnitten, selbst bei in sich relativ geschlossenen wie Teilen zyklischer Werke, kamen oft Abweichungen vor: beispielsweise hatten verschieden lange Ausschnitte ein annähernd gleich umfangreiches Motivinventar und umgekehrt.¹⁰

Ähnliche Erscheinungen wie die oben gezeigten wurden an literarischen Texten festgestellt, für deren Wortwiederholungsstruktur das verallgemeinerte Zipf-Mandelbrotsche Gesetz galt. Deshalb kam es zu der Hypothese, daß die Organisation der Rekurrenz von F-Motiven im melodischen "Schnitt" musikalischer Texte auch für Texte verschiedenster Stile allgemeingültigen Charakter besitzt, daß zwischen ihr und dem Text als Ganzem ein Zusammenhang besteht und daß für sie das Zipf-Mandelbrotsche Gesetz gilt.

3.

Betrachten wir zunächst das jeweilige Motivinventar der untersuchten Texte und überprüfen, inwieweit für vollständige Texte und Textausschnitte das tatsächliche Motivinventar mit der theoretischen Prognose gemäß dem Ausdruck (4), d.h. der Folgerung aus dem verallgemeinerten Zipf-Mandelbrotschen Gesetz, übereinstimmt.

In der Abb. 4 sind die Graphen der Funktion (4) für $p_{\max}=0,04$ und $p_{\max}=0,20$ angegeben, also für den Bereich, in dem die Werte $p_{\max}$ für die untersuchten Texte und Textausschnitte liegen. (Als Textausschnitte wurden einzelne Teile der untersuchten zyklischen Werke, also der Sonaten, der Präludien und Fugen etc., herangezogen.) Als schwarze Punkte sind die jeweiligen Motivinventare der vollständigen Texte eingezeichnet, als kleine Kreise die der Textausschnitte. Wie aus der Abbildung ersichtlich wird, liegen nahezu alle schwarzen Punkte *innerhalb* des Streifens zwischen den Kurven I und II, während die kleinen Kreise größtenteils *außerhalb* dieses Streifens gelegen sind. Die Distribution der kleinen Kreise läßt nur schwer den Schluß auf irgendeine Gesetzmäßigkeit zu. Somit zeigt die Abb. 4, daß das Motivinventar von Textausschnitten im allgemeinen nicht mithilfe des verallgemeinerten Zipf-Mandelbrotschen Gesetzes in der Form der Beziehung (4) beschrieben werden kann, wohingegen das Motivinventar von vollständigen Texten durch dieses Gesetz hinreichend genau beschrieben wird.

In weiteren Einzelheiten sind die Beziehungen zwischen dem "theoretisch" bestimmten und dem "experimentell" ermittelten Motivinventar der vollständigen Texte in der Tabelle 1 (Spalten 1 - 7) dargestellt, in der für jeden Text die Werte für seine

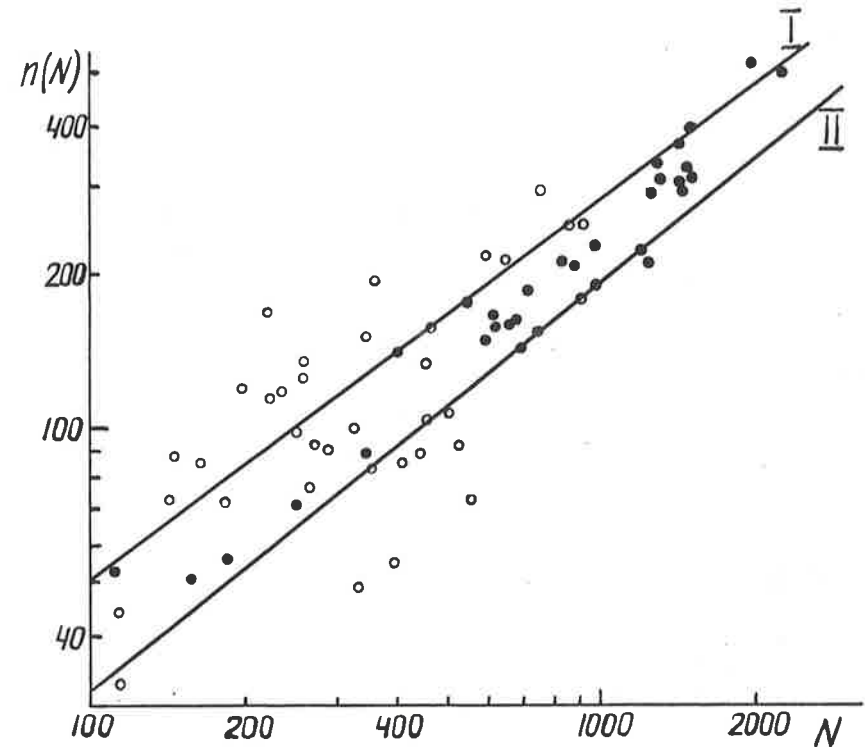


Abb.4

Länge N_0 , für die maximale Vorkommenshäufigkeit eines F-Motivs $p_{\max}$, für das tatsächliche Motivinventar $\underline{n}$ und für dessen Prognose gemäß (4), $\underline{n}^*$, angegeben sind. Wie man dort erkennen kann, sind bei den meisten Texten die Abweichungen von $\underline{n}$ gegenüber der Prognose $\underline{n}^*$ nicht größer als $\pm 20\%$, und bei einer Reihe von Texten (so bei den Sonaten Nr. 1 und Nr. 9 von Scarlatti (6. und 8.), bei der Mozart-Sonate (15.), bei der Fuge von Mjaskovskij (28.) etc.) stimmt das tatsächliche Motivinventar mit seiner Prognose praktisch überein.

Demnach weisen Textausschnitt und vollständiger Text schon aufgrund dieser einen Charakteristik der Wiederholungsstruktur von F-Motiven, dem Motivinventar, einen wesentlichen Unterschied

auf. Der Vergleich zwischen dem Motivinventar vollständiger Texte und dem von Textausschnitten mit den entsprechenden Prognosen gemäß (4) zeigt, daß für die Wiederholungsstruktur von F-Motiven bei Textausschnitten insgesamt das verallgemeinerte Zipf-Mandelbrot'sche Gesetz nicht gilt, wohingegen die Wiederholungsstruktur von F-Motiven im vollständigen Text diesem Gesetz gehorcht. Diese Feststellung erlaubt es, mithilfe des Gesetzes (3) zu einer detaillierteren Analyse der Organisation der Rekurrenz von F-Motiven in vollständigen musikalischen Texten überzugehen.

Wir werden die Häufigkeitsstrukturen von Texten betrachten, und wir werden den tatsächlichen Verlauf der Häufigkeitsabnahme von F-Motiven im Text mit seiner Beschreibung gemäß (3) vergleichen. In der Abb. 5.1 sind die experimentell und theoretisch ermittelten Graphen der Häufigkeitsstrukturen einiger Texte angegeben:

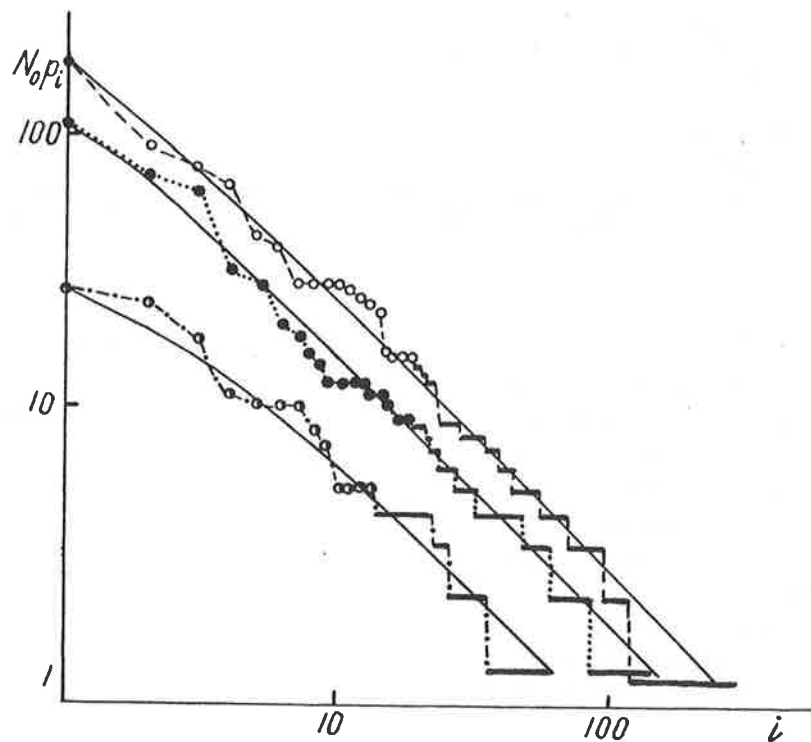


Abb. 5.1.1

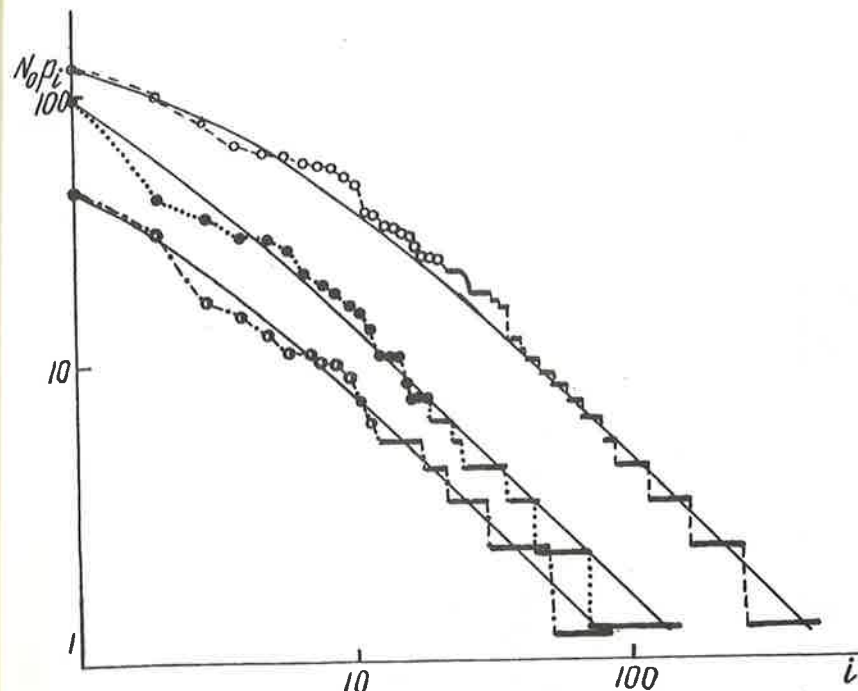


Abb. 5.1.2

Abb. 5.1: Graphische Darstellung der Häufigkeitsstrukturen vollständiger musikalischer Texte.

Auf der Ordinatenachse sind die absoluten Vorkommenshäufigkeiten der F-Motive im Text angetragen, auf der Abszissenachse die Ränge (Ordnungsnummern) der Häufigkeiten in einer nach abnehmenden Werten geordneten Reihenfolge. Auf beiden Achsen ist der Maßstab logarithmisch. Die theoretischen Kurven (gemäß Formel (3) konstruiert) erscheinen in Form einer durchgezogenen Linie.

Abb. 5.1.1:

- a) J.S. Bach. Präludium und Fuge Nr. 22 aus dem Wohltemperierten Klavier, 2. Teil BWV 891 (O--O--O)
- b) J.A. Levitin. Sonatine für Flöte solo (O...O...O)
- c) D. Scarlatti. Sonate für Klavier Nr. 1 (Edition Peters) (O--●--O)

Abb. 5.1.2:

- d) F. Chopin. Sonate Nr. 3 h-moll op. 58 (O--O--O)
- e) J.S. Bach. Präludium und Fuge Nr. 2 aus dem Wohltemperierten Klavier, 2. Teil BWV 871 (●...●...●)
- f) L. van Beethoven. Sonatine für Klavier F-Dur (O--●--O)

Wie die Abbildung zeigt, stimmen die experimentell ermittelten Kurven mit ihren theoretischen Beschreibungen gemäß (3) gut überein, sogar für einen relativ kurzen Text wie die Sonate Nr. 1 von Scarlatti (mit einer Länge von nur 250 F-Motiven). Gleichzeitig sind dabei die Texte ihrem Stil, ihrer Kompositionsform und ihrem satztechnischen Typ (homophone und polyphone Texte) nach unterschiedlich.

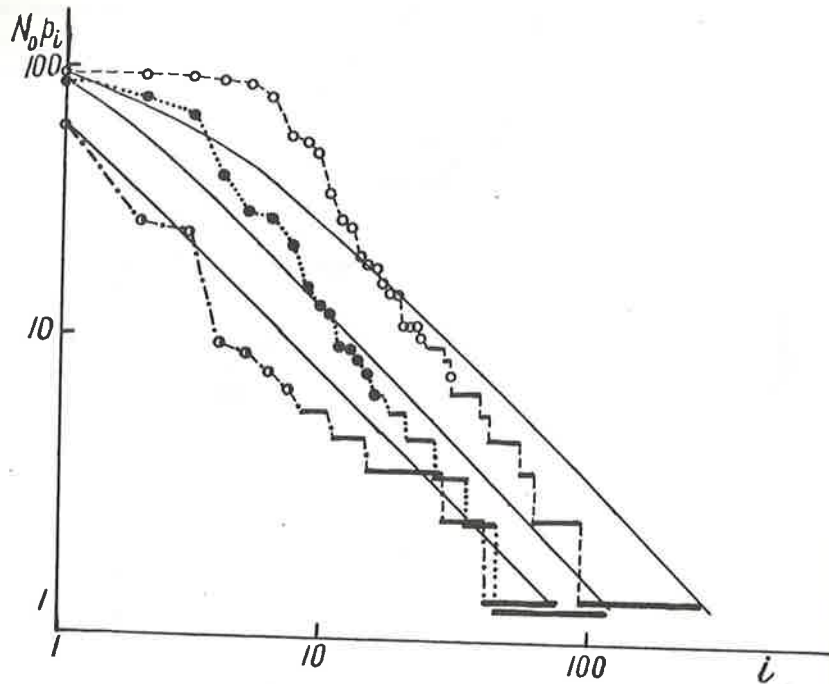


Abb. 5.2.1

Abb. 5.2: Graphische Darstellung der Häufigkeitsstrukturen von Textausschnitten (Teilen zyklischer Werke)

Die Benennungen der Achsen und der Maßstab sind dieselben wie in Abb. 5.1. Die theoretischen Kurven erscheinen in Form einer durchgezogenen Linie.

- Abb. 5.2.1: g) J.S. Bach. Contrapunctus 8 aus der Kunst der Fuge (O--O--O)
 h) J.S. Bach. Präludium aus BWV 891 (●...●...●)
 i) Ju.A. Levitin. Sonatine für Flöte solo, 1. Satz (●--●)

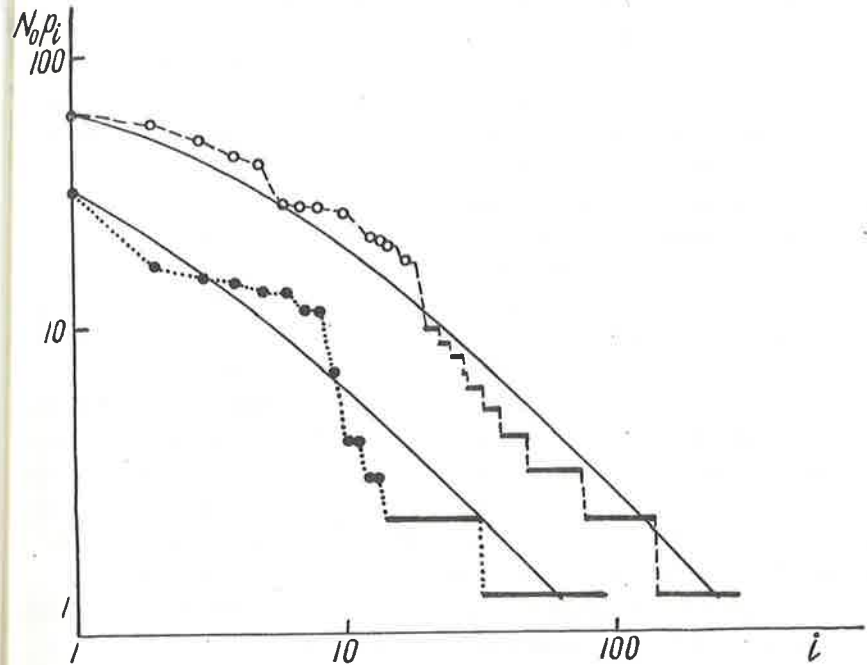


Abb. 5.2.2

Abb. 5.2.2:

j) F. Chopin. Sonate Nr. 3, 1. und 2. Satz (O--O)

k) D.D. Sostakovič. Präludium aus op. 87 Nr. 4 (●...●)

Zum Vergleich sind in der Abb. 5.2 die Graphen der Häufigkeitsstrukturen von Textausschnitten angegeben, also von einzelnen Teilen zyklischer Werke (darunter von solchen, für die die Häufigkeitsstrukturen der vollständigen Texte in der Abb. 5.1 dargestellt sind).

Die theoretischen Werte der Häufigkeiten der F-Motive sind nach der Formel (3) berechnet für N_0 gleich der Länge des Textausschnittes und $p_{\max}$ gleich der Häufigkeit seines häufigsten F-Motives.

Wie aus der Abb. 5.2 klar ersichtlich wird, weichen die experimentell ermittelten Häufigkeitskurven ganz erheblich von den zugehörigen theoretischen Kurven ab, und zwar im Bereich kleiner Häufigkeiten

nach unten, im Bereich großer Häufigkeiten gewöhnlich nach oben, und bei einigen Graphen zeigen sich auch im Bereich mittlerer Häufigkeiten Abweichungen ("Durchbiegungen") der experimentellen von der theoretischen Kurve (so z.B. bei dem Satz der Levitin-Sonatine).

Dabei ist es sehr wichtig, darauf hinzuweisen, daß die Längen der Textausschnitte, deren Häufigkeitsstrukturen in Abb. 5.2 (1+2) dargestellt sind, durchaus vergleichbar sind mit den Längen der vollständigen Texte, die Abb. 5.1 (1+2) zugrundeliegen. So beträgt beispielsweise die Länge des Ausschnittes aus Bachs Variationen-Zyklus "Die Kunst der Fuge" (Contrapunctus 8, vgl. Abb. 5.2.1) 135 F-Motive und ist damit annähernd gleich der Länge von Präludium und Fuge Nr. 22 aus dem Wohltemperierten Klavier, 2. Teil (BWV 891, vgl. Abb. 5.1.1). Beim Vergleich von Abb. 5.1.1 und Abb. 5.2.1 läßt sich jedoch ein deutlicher Unterschied zwischen der Häufigkeitsstruktur der F-Motive des vollständigen Textes und der des Textausschnittes erkennen. Ein ähnliches Bild ergibt sich, wenn man den Ausschnitt aus Präludium und Fuge Nr. 22 (BWV 891, vgl. Abb. 5.2.1) dem annähernd gleichlangen vollständigen Text von Präludium und Fuge Nr. 2 aus dem Wohltemperierten Klavier, 2. Teil (BWV 871, vgl. Abb. 5.1.2) gegenüberstellt, annähernd gleichlange vollständige Texte von Chopin Textausschnitten etc. In allen diesen Fällen lassen sich bei den Textausschnitten erhebliche Abweichungen vom Zipf-Mandelbrot'schen Gesetz feststellen: Der vollständige Text erfüllt dieses Gesetz wesentlich besser als der annähernd gleichlange Textausschnitt, selbst bei Werken desselben Komponisten. Ein ähnliches Bild wie in Abb. 5.1 und 5.2 läßt sich auch für die anderen untersuchten vollständigen Texte und Textausschnitte feststellen.

All diese Beobachtungen lassen die Aussage zu, daß in den untersuchten musikalischen Texten aus einem breiten Spektrum von Stilen die Organisation der Rekurrenz von F-Motiven in der Tat allgemeingültigen Abhängigkeitsverhältnissen unterliegt. Sie kann mithilfe des Zipf-Mandelbrot'schen Gesetzes beschrieben werden, doch darüber hinaus steht die Erfüllung dieses Gesetzes in unmittelbarem und wesentlichem Zusammenhang mit dem musikalischen Werk als Ganzem, mit seiner Vollständigkeit und Abgeschlossenheit. Infolgedessen steht die Wiederholungsstruktur von F-Motiven im Text für alle untersuchten musikalischen Texte in einem gleichartigen Zusammenhang mit der vollen Textlänge.

4.

Die letztgenannte Schlußfolgerung hat eine wesentliche Konsequenz: Wenn die Wiederholungsstruktur von F-Motiven in einem für alle untersuchten Texte gleichartigen Zusammenhang steht, dann müssen in annähernd gleichlangen Texten (mit gleichzeitig annähernd gleichen Werten für die Häufigkeit ihres jeweiligen häufigsten F-Motivs $p_{\max}$) auch ähnliche bzw. annähernd gleiche Wiederholungsstrukturen von F-Motiven vorliegen. Demgegenüber müssen sie sich in Texten mit erheblichen Längenunterschieden voneinander unterscheiden. Dabei spielt es keine Rolle, ob die Texte ihrem Stil nach grundsätzlich verschieden oder annähernd gleich sind.

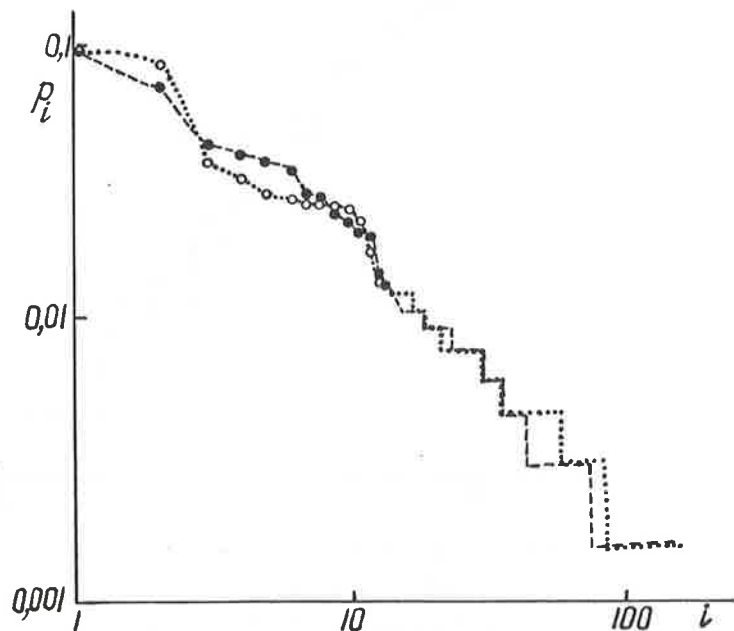


Abb. 6.1: Graphen der Häufigkeitsstrukturen musikalischer Texte mit annähernd gleicher Länge (bei annähernd gleichen Werten für $p_{\max}$).

- a) D.B. Kabalevskij. Rondo für Klavier op. 59 (Länge: 625 F-Motive) (O...O)
- b) L. van Beethoven. Rondo C-Dur für Klavier op. 51 Nr. 1 (Länge: 624 F-Motive) (●--●)

Wie die Analyse zeigt, liegt diese maßgebende Erscheinung tatsächlich vor. So sind in der Abb. 6.1 beispielsweise die Graphen zweier Texte angegeben, deren Länge und ebenso deren Werte für $p_{\max}$ praktisch gleich sind.

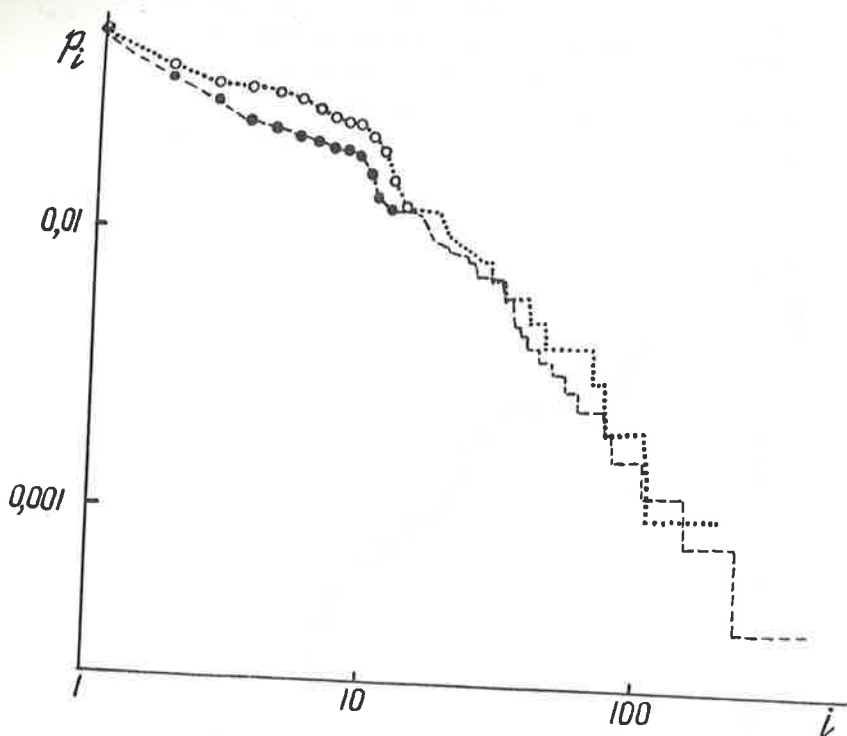


Abb. 6.2: Graphen der Häufigkeitsstrukturen musikalischer Texte mit unterschiedlicher Länge (bei annähernd gleichen Werten für $p_{\max}$).

- c) F. Chopin. Phantasie f-moll op. 49 (Länge: 987 F-Motive) (O...O)
 d) F. Chopin. Sonate Nr. 3 h-moll op. 58 (Länge: 2364 F-Motive) (●--●)

Wie aus der Abbildung ersichtlich wird, ist trotz aller grundlegenden stilistischen Unterschiedlichkeit der Texte erwartungsgemäß der allgemeine Verlauf der Häufigkeitskurve für beide Rondos annähernd gleich, es zeigt sich weder ein systematisch höherer noch ein systematisch niedrigerer Verlauf einer Häufigkeitskurve im Vergleich zur anderen.

Im Gegensatz dazu sind in der Abb. 6.2 die Graphen der Häufigkeitsstrukturen von Texten angegeben, die zwar auch annähernd gleiche Werte für $p_{\max}$ haben, sich aber bezüglich ihrer Länge im Verhältnis 1:2,5 unterscheiden.

In diesem Falle ist, wie man klar erkennen kann, der Verlauf der Häufigkeitsabnahme in den beiden Texten unterschiedlich: Die Kurve für die Phantasie (den kürzeren Text!) verläuft insgesamt höher und fällt flacher als die Kurve für die Sonate. Ähnliche Verhältnisse, wie sie in Abb. 6.1 und 6.2 dargestellt sind, lassen sich auch für einige andere der untersuchten Texte aufzeigen.

5.

Die angeführten Beispiele werfen aber folgendes Problem auf: Wenn bei annähernder Gleichheit der Werte für $p_{\max}$ die Häufigkeitskurve für den längeren Text insgesamt unter der Kurve für den längeren Text verläuft, dann muß der längere Text relativ weniger häufig vorkommende F-Motive als der kürzere Text enthalten und andererseits relativ mehr selten, insbesondere nur einmal vorkommende. Infolgedessen müssen neue (vom Textanfang an erstmalig vorkommende) F-Motive im längeren Text durchschnittlich häufiger auftreten als im kürzeren. Und folglich muß, beim Vergleich von gleichlangen Textausschnitten aus einem längeren und einem kürzeren Text, das Motivinventar des Ausschnittes aus dem längeren durchschnittlich größer sein. (Der Zusatz "durchschnittlich" soll auf die Notwendigkeit hinweisen, die Ungleichmäßigkeiten im Anwachsen des Motivinventars im Text zu berücksichtigen, die durch die Eigenart der jeweiligen musikalischen Form bedingt sind). Mit anderen Worten: Der längere Text muß an F-Motiven (in der Terminologie der vorliegenden Arbeit: bezüglich seines Motivinventars) stärker gesättigt sein als der kürzere.

Diese Schlußfolgerung erscheint, wenn man die grundlegenden stilistischen Unterschiede der untersuchten Texte bedenkt (vgl. Tabelle 1), bei weitem nicht offenkundig. A priori vermutet man, daß bei einem musikalischen Text der "motivische Sättigungsgrad" in erster Linie durch den Stil und die Entstehungszeit des betreffenden Werkes bedingt ist. Es mutet wahrscheinlich an, daß die Werke eines einzelnen Komponisten (jedenfalls die größeren) alle unge-

fähr denselben motivischen Sättigungsgrad besitzen, zumindest auf der Ebene der F-Motive. Gerade deshalb ist es notwendig, den Zusammenhang zwischen dieser Charakteristik und der Textlänge experimentell zu überprüfen.

In der Abb. 7 sind die Graphen für die Akkumulation des Motivinventars, d.h. der Anzahl der distinkten F-Motive im Text, für einige verschieden lange Texte angegeben, von denen drei vom selben Komponisten (Chopin) stammen:

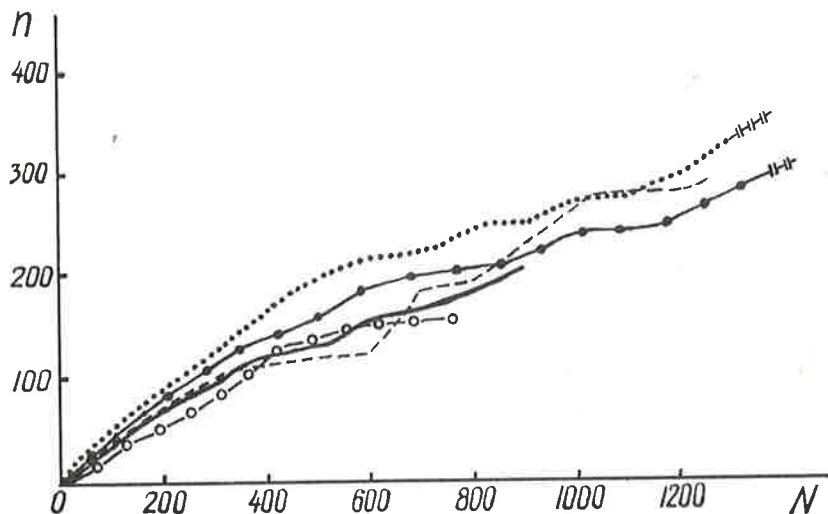


Abb. 7: Graphische Darstellung des Anwachsens des Motivinventars im musikalischen Text.

(Der weitere Verlauf der Kurven a) und c) über den Rand der Abbildung hinaus ist durch ---|---| angedeutet).

- a) F. Chopin. Sonate Nr. 3 h-moll op. 58 (•••••)
- b) W.A. Mozart. Sonate für Klavier Nr. 10 (---)
- c) F. Chopin. Sonate Nr. 2 b-moll op. 35 (•••••)
- d) F. Chopin. Ballade Nr. 1 g-moll op. 23 (—+—+—+—+—)
- e) Ju.A. Levitin. Sonatine für Flöte solo (o-o-o-o)

Die Abbildung zeigt, daß sich ungeachtet der auffälligen Unregel-

mäßigkeit im Anwachsen des Motivinventars im Text¹¹ deutlich die Tendenz beobachten läßt, daß der längere Text motivisch stärker gesättigt ist. Die Kurve des Anwachsens des Motivinventars im längeren Text verläuft im allgemeinen über der entsprechenden Kurve des kürzeren Textes. Dies ist besonders auffällig zu erkennen bei den drei Chopin-Texten: der Ballade Nr. 1 (Kurve d); Länge: ca. 900 F-Motive), der Sonate Nr. 2 (Kurve c); Länge: ca. 1500 F-Motive) und der Sonate Nr. 3 (Kurve a); Länge ca. 2400 F-Motive). Dort verläuft nämlich die Kurve des Anwachsens des Motivinventars für die Ballade insgesamt unter der Kurve für die Sonate Nr. 2, und diese wiederum verläuft unter der Kurve für die Sonate Nr. 3.

Somit wird für die untersuchten musikalischen Texte das verallgemeinerte Zipf-Mandelbrotsche Gesetz "im großen und ganzen" erfüllt (man denke an den übereinstimmenden Verlauf der experimentellen und theoretischen Häufigkeitskurven und die Übereinstimmungen mit der Beziehung (4) bezüglich des Motivinventars), und darüber hinaus ergibt sich aus diesem Gesetz eine wesentliche Folgerung in Form einer deutlichen Tendenz: Je länger ein Text ist, desto höher ist sein motivischer Sättigungsgrad.

6.

Die experimentellen Graphen der Häufigkeitsstrukturen in den Abb. 5. zeigen noch einen weiteren wichtigen Aspekt der Organisation der Rekurrenz von F-Motiven in den untersuchten Texten. Bei jedem der Graphen lassen sich in seinem rechten Teil charakteristische "Stufen" erkennen, die mit dem Vorhandensein umfangreicher Gruppen von gleich häufig vorkommenden F-Motiven mit kleiner Häufigkeit in Zusammenhang stehen. Wie aus den Abbildungen ersichtlich, wird der allgemeine Verlauf der Häufigkeitsabnahme von F-Motiven durch das verallgemeinerte Zipf-Mandelbrotsche Gesetz (3) im allgemeinen gut beschrieben. Jedoch wird die theoretische Kurve von den genannten "Stufen" vielfach geschnitten, und teilweise liegen diese auch unterhalb der Kurve, ohne sie überhaupt zu berühren. Dies ist ein Anzeichen dafür, daß die reale Wiederholungsstruktur von F-Motiven in musikalischen Texten in bestimmter Weise vom Gesetz (3) abweicht. In Anbetracht dessen, daß der

Bereich kleiner Häufigkeiten einen bedeutenden Teil der Häufigkeitsstruktur ausmacht, war es also von Interesse nachzuprüfen, wie groß diese Abweichungen sind und wie exakt die Zahl der m -mal vorkommenden F-Motive durch den Ausdruck (5), d.h. durch die Folgerung aus dem Gesetz (3) beschrieben wird. Die entsprechenden Daten sind in der Tabelle 1 (Spalten 8 - 16) aufgelistet. Die tatsächliche Anzahl der m -mal vorkommenden F-Motive ist dort mit n_m bezeichnet, die zugehörige Prognose gemäß (5) mit n_m^* .

Wie man erkennen kann, sind die Abweichungen von n_m gegenüber der Prognose n_m^* im Durchschnitt erheblich größer als die Abweichungen beim Motivinventar. Da aber der allgemeine Verlauf der Häufigkeitsabnahme und das Motivinventar von Texten mit dem verallgemeinerten Zipf-Mandelbrotschen Gesetz in Übereinstimmung stehen, läßt sich vermuten, daß im Bereich der kleinen Häufigkeiten gegenseitige Kompensationsmechanismen wirksam sind, bei denen ein Defizit an Elementen mit der einen Häufigkeit durch einen Überschuß an Elementen mit einer anderen sozusagen wieder ausgeglichen wird. In der Tat kann man anhand der Daten der Tabelle 1 feststellen, daß in einem Text, bei dem beispielsweise die Anzahl der einmal vorkommenden F-Motive kleiner ist als ihre Prognose gemäß (5), die Anzahl der zweimal oder dreimal vorkommenden F-Motive entsprechend größer ist als die zugehörige Prognose. Ähnliches gilt auch in anderen Fällen. Offensichtlich gewährleisten eben derartige "Umgruppierungen" die insgesamt gute Übereinstimmung der experimentellen Häufigkeitskurven mit dem verallgemeinerten Zipf-Mandelbrotschen Gesetz.

Anhand der Daten der Tabelle 1 läßt sich auch die allgemeine Ursache angeben, durch die wahrscheinlich diese "Umgruppierungen" zustandekommen. Wenn man die Zahl der m -mal im Text vorkommenden F-Motive mit der zugehörigen Prognose gemäß (5) vergleicht, kann man erkennen, daß die Zahl der *zweimal vorkommenden* F-Motive in der Regel *größer* ist als ihre theoretische Prognose. In drei Texten (in der Sonate von Tartini, in der Klaviersonate und in der "Kleinen Nachtmusik" von Mozart) ist die Zahl der zweimal vorkommenden F-Motive sogar *größer als die der einmal vorkommenden*. Im Durchschnitt des ganzen Textkorpus ist die Abweichung der Zahl der zweimal vorkommenden F-Motive gegenüber der Prognose gemäß (5) ungefähr dreimal so groß wie die entsprechenden Abweichungen

für die einmal bzw. dreimal vorkommenden F-Motive (bei letzteren dem Betrage nach). All diese Beobachtungen lassen die Annahme zu, daß in den untersuchten Texten ein eigenartiges "Prinzip des zweimaligen Vorkommens eines F-Motivs" wirksam ist, die Tendenz, ein einmal vorgekommenes F-Motiv zu *wiederholen*. Dieses Prinzip ist offenbar spezifisch für musikalische Texte (als "Prinzip des zweimaligen Vorkommens eines Elementes"), denn bei literarischen Texten, die bezüglich der Geltung des Gesetzes (3) untersucht worden sind, weicht, wie sich den Daten der Arbeit von ORLOV (1975) entnehmen läßt, die Zahl der zweimal vorkommenden Wörter nur unwesentlich von (5) ab, und im Durchschnitt sind diese Abweichungen annähernd gleich Null. Überdies kann die übergroße Zahl der zweimal vorkommenden F-Motive in Zusammenhang gebracht werden mit der für die Musik charakteristischen Tendenz, ein Element *unmittelbar* oder *fast unmittelbar nach seinem ersten Vorkommen zu wiederholen*, so wie es bei der Wiederholung eines Themas nach seiner ersten Formulierung, der Wiederholung der Exposition einer Sonate oder Symphonie, der Motiv-Wiederholung bei der Formulierung eines Themas etc. geschieht. Offenbar ist dasselbe Prinzip auch auf der Ebene der F-Motive (vgl. auch BORODA 1976) gültig. Unter diesen Umständen kann das Defizit an einmal vorkommenden F-Motiven in einer Reihe von Texten darauf zurückzuführen sein, daß die Tendenz zur Elementwiederholung derartig stark in Erscheinung tritt, daß sie teilweise auf "potentiell nur einmal vorkommende" F-Motive übergreift. Es ist durchaus wahrscheinlich, daß das verallgemeinerte Zipf-Mandelbrotsche Gesetz und das "Prinzip des zweimaligen Vorkommens eines F-Motivs" für die untersuchten Texte gewissermaßen *fundamentale Gegebenheiten* darstellen. Die "Kompensationsercheinungen" und generell die erheblichen Abweichungen im Bereich kleiner Häufigkeiten kommen dann zustande infolge der Kollision dieser beiden fundamentalen Gegebenheiten miteinander.

Fassen wir also unsere Ergebnisse zusammen. Wir haben anhand von musikalischem Textmaterial aus den Stilen des 18. - 20. Jahrhunderts die Organisation der Rekurrenz von kleinen melodischen Elementen, den F-Motiven, im musikalischen Text untersucht. Diese Analyse ergab

folgendes:

1. Die Wiederholungsstruktur von F-Motiven im Text steht in Zusammenhang mit der Textlänge N_0 (der Summe aller gebrauchten F-Motive) und mit der Vorkommenshäufigkeit seines häufigsten F-Motivs $p_{\max}$. Diese Struktur läßt sich mithilfe einer Form des in der quantitativen Linguistik bekannten Zipf-Mandelbrotschen Gesetzes beschreiben. Dabei ist es wesentlich, den Text als Ganzes zu betrachten, denn für Textausschnitte, selbst für in sich relativ geschlossene wie Teile von Werken in zyklischer Form, wird dieses Gesetz nicht erfüllt.

2. Je länger ein musikalischer Text ist, desto umfangreicher ist, bei gleichen Werten für $p_{\max}$, sein Motivinventar (die Anzahl der distinkten F-Motive) und desto höher ist sein motivischer Sättigungsgrad. Demnach läßt sich bei einem Ausschnitt aus einem langen Text ein umfangreicheres Motivinventar feststellen als bei einem gleichlangen Ausschnitt aus einem kürzeren Text.

3. In jedem Text gibt es eine beträchtliche Anzahl selten vorkommender F-Motive. Die Zahl der m -mal vorkommenden F-Motive wächst mit abnehmendem m . Die Analyse des Bereichs kleiner Häufigkeiten ($m = 1, 2, 3$) ließ erkennen, daß in den Texten ein "Prinzip des zweimaligen Vorkommens eines F-Motivs" wirksam ist, das mit der für die Musik charakteristischen Tendenz in Zusammenhang steht, ein einmal vorgekommenes Element zu wiederholen, es ein zweites Mal vorkommen zu lassen. Diese Tendenz, die bei der ansonsten im allgemeinen dem Zipf-Mandelbrotschen Gesetz entsprechenden Wiederholungsstruktur von F-Motiven gleichzeitig in Erscheinung tritt, ruft offenbar die an vielen Texten festgestellten "Kompensationserscheinungen" hervor, bei denen der Überschuß an zweimal vorkommenden F-Motiven durch ein Defizit an dreimal, und in einigen Fällen auch an einmal vorkommenden F-Motiven ausgeglichen wird.

Übersetzt aus dem Russischen von Hans-Joachim KEMPER

Anmerkungen:

1 Die Untersuchungen zu dieser Arbeit wurden am Lehrstuhl für Ästhetik und Kunstwissenschaft des Staatlichen V. Saradžišvili-Konservatoriums Tbilisi (Kafedra estetiki i iskusstvovedenija Tbilisskoj gosudarstvennoj konservatorii im. V. Saradžišvili) durchgeführt. Der Aufsatz, dessen überarbeitete und ergänzte Fassung hier in Übersetzung erscheint, wurde in russischer Sprache unter dem Titel "Častotnye struktury muzykal'nych tekstov" in dem Sammelband: Sbornik statej posvjaščennyj 60-letiju Velikoj Oktjabr'skoj socialističeskoj revoljucii (Sammelband aus Anlaß der 60-Jahrfeier der Oktoberrevolution). (ed. A. SAVERZAŠVILI et al.). Tbilisi: Mecniereba, 1977 veröffentlicht. Der Autor dankt Ju.K. ORLOV, R.Ch. ZARIPOV und E.M. DUMANIS für ihre freundlichen Hinweise bei der Durchsicht der vorliegenden Arbeit.

2 Gemeint sind die Organisationsprinzipien für einen Text in seiner Gesamtheit, von Anfang bis Ende. Die Gesetzmäßigkeiten der Wiederholung und Alternation kleiner Elemente in weniger umfangreichen Teileinheiten, insbesondere in thematischen Abschnitten, sind in der Musikwissenschaft auch in Einzelheiten hinreichend bekannt (vgl. MAZEL', CUKKERMANN 1967).

3 Unter der Textlänge versteht man die Zahl der im Text gebrauchten Wortformen.

4 In Anlehnung an ORLOV (1975) werden wir den Ausdruck (3) im folgenden auch als verallgemeinertes Zipf-Mandelbrotsches Gesetz bezeichnen.

5 Die Namen russischer bzw. sowjetischer Autoren (Schriftsteller und Komponisten) werden hier nicht in der sonst üblichen Weise transkribiert, sondern sie erscheinen, ebenso wie die von Verfassern wissenschaftlicher Arbeiten, in wissenschaftlicher Transliteration (also z.B. Šostakovič - nicht Schostakowitsch). (Anm. d. Übers.).

6 Eine detailliertere und enger gefaßte Definition für a) - d) findet sich in unseren erwähnten Arbeiten (BORODA 1973, 1977).

7 In polyphonen Texten wurde zunächst jede einzelne Stimme *s e - p a r a t* in F-Motive segmentiert, und daraufhin wurden die jeweiligen Summen der 1., 2. etc. Stimme in einer einzigen Summe vereinheitlichend zusammengefaßt, anhand derer dann die unten aufgeführten Charakteristika für den Text bestimmt wurden.

8 Der im russ. Original gebrauchte Terminus, den wir hier mit "Motivinventar" übersetzen, lautet "intonacionnyj zapas" (wörtlich: intonatorisches Inventar). In der russ. musikwissenschaftlichen Terminologie hat "intonacija" (wörtlich: Intonation) einerseits dieselben Bedeutungen wie der deutsche Terminus "Intonation", andererseits wird es aber auch in bestimmten, im hier gegebenen Zusammenhang vorliegenden, Fällen synonym für "Motiv" gebraucht. Da dies in der deutschen Terminologie nicht möglich ist, mußte auf die Übersetzung "Motivinventar" zurückgegriffen werden, auch wenn dabei die terminologischen Unterschiede zwischen "intonacija" und "motiv" bzw. zwischen "F-Motiv" und "Motiv" in der gängigen Definition nicht deutlich gemacht werden. Darüber hinaus ist der Gebrauch von "intonacija" hier insbesondere auch dadurch motiviert, daß in metaphorischer Bedeutung "intonacija" als "musikalisches Wort" (muzykal'noe slovo) verstanden werden kann. Somit ist die hier aufzuzeigende Analogie zu literarischen Texten, an denen das Lexeminventar (slovarnyj zapas) untersucht wird, durch den Terminus "intonacionnyj zapas" für "Motivinventar" im Sinne von "Inventar an musikalischen Wörtern" zusätzlich verdeutlicht. Zu den Bedeutungen von "motiv" und "intonacija" in der russ. Terminologie, besonders auch zum Verhältnis von "intonacija" als musikalischem Wort zum Wort als lexikalischer Einheit vgl. das neueste sowjetische musikwissenschaftliche Standardwerk: Muzykal'naja enciklopedija. (ed. Ju.V. KELDYŠ). Bd. 1, Moskva, 1973. "motiv": Bd. 3, 1976, Spalte 696-698; "intonacija": Bd. 2, 1974, Sp. 550-557, bes. Sp. 553-554. (Anm. d. Übers.).

9 Dabei war natürlich die tongetreue Wiederholung eines F-Motivs als eine Sequenzbildung mit Null-Intervall mitverstanden.

10 Die Textlänge, das Motivinventar und die Vorkommenshäufigkeiten von F-Motiven wurden für Textausschnitte in der gleichen Weise bestimmt wie für vollständige Texte.

11 Zu einigen Erscheinungsformen und Ursachen dieser Unregelmäßigkeit vgl. BORODA (1976).

Erläuterungen zu Tabelle1:

N_0 - Textlänge; p_{max} - relative Häufigkeit des häufigsten F-Motivs im Text; n - tatsächliches Motivinventar des Textes; n^* - Prognose für das Motivinventar nach der Formel (4); δ - relative Abweichung von n gegenüber n^* ; n_m ($m = 1, 2, 3$) - tatsächliche Anzahl der einmal, zweimal bzw. dreimal im Text vorkommenden F-Motive; n_m^* - theoretische Prognose für die Anzahl der m -mal im Text vorkommenden F-Motive nach der Formel (5); δ_m - relative Abweichung von n_m gegenüber n_m^* . Ebenso wird in der Tabelle für jeden Text der Wert des Verhältnisses n_1/n der Anzahl der einmal vorkommenden F-Motive zum Motivinventar angegeben.

Die Korrespondenzen der hier angegebenen Numerierung zu der in neuerer Zeit üblich gewordenen Numerierung der Werke D. Scarlattis nach Kirkpatrick konnten für die in Frage kommenden Sonaten mangels Material nicht ermittelt werden. (Anm. d. Übers.)

Tabelle 1

1	2	3	4	5	6
lfd. Nr.	Materialquelle	N _O	P _{max}	n	n*
1.	J.S. Bach (1685-1750). Präludium und Fuge Nr. 13 BWV 858; aus dem Wohltemperierten Klavier, 1. Teil (abgek.: WTK 1).	615	0.0910	186	153.0
2.	ders. Präludium und Fuge Nr. 20 BWV 865; WTK 1.	1422	0.1958	296	252.0
3.	ders. Präludium und Fuge Nr. 2 BWV 871; WTK 2.	671	0.1430	168	145.0
4.	ders. Präludium und Fuge Nr. 22 BWV 891; WTK 2.	1430	0.1350	310	272.0
5.	G.F. Händel (1685-1756). Fuge für Orgel Nr. 2 G-Dur.	765	0.1070	201	173.0
6.	D. Scarlatti (1685-1757). Sonate Nr. 1 (Zählung der Edition Peters).	250	0.1190	72	72.0
7.	ders. Sonate Nr. 5.	182	0.0659	60	67.0
8.	ders. Sonate Nr. 9. (Ausgabe von B. Holdenweiser).	156	0.1090	53	52.0
9.	ders. Sonate Nr. 13. (Ausgabe von L. Nikolaev).	568	0.0986	190	139.0
10.	ders. Sonate Nr. 25 (Edition Peters).	107	0.0655	51	47.2
11.	G. Tartini (1692-1770). Sonate g-moll für Violine und Klavier.	828	0.0652	218	204.0
12.	J. Haydn (1732-1809). Symphonie Nr. 45 fis-moll "Abschieds-Symphonie".	1304	0.0437	340	317.0
13.	W.A. Mozart (1756-1791). "Eine kleine Nachtmusik" KV 525.	1260	0.1143	205	252.0
14.	ders. Fuge g-moll für Klavier.	731	0.1370	199	157.0
15.	ders. Sonate für Klavier Nr. 10.	1249	0.0509	296	297.0
16.	J.G. Albrechtsberger (1736-1809). Fuge c-moll.	693	0.2160	168	138.0
17.	L. van Beethoven (1770-1827). Sonate für Klavier Nr. 19.	584	0.0565	151	163.0

7	8	9	10	11	12	13	14	15	16	17
δ	n_1	n_1^*	δ_1	n_2	n_2^*	δ_2	n_3	n_3^*	δ_3	$\frac{n_1}{n}$
21.6	105	93.0	12.9	31	31.0	0.0	12	15.5	-22.6	0.565
17.45	192	126.0	58.4	39	42.0	-7.2	16	21.0	-23.8	0.648
15.8	96	72.5	32.4	26	24.2	8.3	9	12.1	-25.8	0.572
13.6	191	136.0	41.3	26	45.0	-42.0	21	22.7	-7.5	0.616
16.2	120	86.5	38.8	28	28.8	-2.8	8	14.0	-42.9	0.598
0.0	37	36.0	2.8	9	12.0	-25.0	3	6.0	-50.0	0.515
-10.5	28	33.5	-16.2	16	11.2	42.9	1	5.6	-46.6	0.467
1.9	20	26.0	-23.1	11	8.7	27.0	8	4.3	84.5	0.378
36.8	100	69.5	43.9	39	23.2	68.2	17	11.6	46.6	0.526
8.1	25	23.6	5.9	15	7.9	90.7	3	3.4	-12.5	0.490
6.9	60	102.0	-41.2	85	34.0	150.0	11	17.0	-35.3	0.276
7.3	164	158.5	3.5	65	52.8	23.1	25	28.5	-6.0	0.483
-18.7	60	126.0	-52.4	61	42.2	44.6	13	21.1	-38.4	0.293
26.7	119	78.5	51.6	33	25.8	16.3	10	12.9	-14.7	0.548
0.0	95	149.0	-36.2	98	49.5	98.0	14	24.8	-43.6	0.321
20.9	111	69.8	70.0	17	23.0	-21.8	8	11.5	-30.8	0.661
-7.4	60	81.5	-26.4	36	27.2	32.4	13	13.6	-4.4	0.400

1	2	3	4	5	6
18.	ders. Rondo C-Dur für Klavier op. 51 Nr. 1.	624	0.0960	162	150.0
19.	ders. Sonatine für Klavier F-Dur.	333	0.1320	85	86.6
20.	R. Schumann (1810-1856). Konzert für Violoncello und Orchester a-moll op. 129.	2023	0.0537	532	428.0
21.	F. Chopin (1810-1849). Ballade Nr. 1 g-moll op. 23.	902	0.0932	206	201.0
22.	ders. Sonate Nr.2 b-moll op.35.	1484	0.0458	330	346.0
23.	ders. Sonate Nr.3 h-moll op.58.	2364	0.0532	504	485.0
24.	ders. Phantasie f-moll op. 49.	987	0.0589	209	239.0
25.	F. Mendelssohn-Bartholdy (1809-1847). Präludium und Fuge für Orgel e-moll.	1393	0.1098	282	276.0
26.	C. Saint-Saëns (1835-1921). Introduction et Rondo capriccioso für Violine und Orchester.	1200	0.0675	222	270.0
27.	A.N. Skrjabin (1872-1915). Valse op. 38.	402	0.0746	155	117.0
28.	N.J. Mjaskovskij (1881-1950). Fuge für Klavier e-moll.	214	0.1635	60	58.5
29.	S.S. Prokof'ev (1891-1953). Sonate für Violine solo op. 115.	1519	0.0745	398	318.0
30.	P. Hindemith (1895-1963). Sonate für Violine solo Nr. 1	1452	0.0584	374	325.0
31.	D.D. Šostakovič (1906-1975). Präludium und Fuge für Klavier op. 87 Nr. 4	983	0.1078	237	208.0
32.	ders. Präludium und Fuge für Klavier op. 87 Nr. 9	677	0.2160	146	135.0
33.	ders. Präludium und Fuge für Klavier op. 87 Nr. 12	1292	0.0844	316	274.0
34.	D.B. Kabalevskij (geb. 1904). Rondo für Klavier op. 59.	625	0.0950	171	150.0
35.	S.M. Taktakišvili (geb. 1900). Rondo für Violine und Klavier.	512	0.1268	103	120.0
36.	Ju.A. Levitin (geb. 1912). Sonatine für Flöte solo.	760	0.1466	159	160.0

7	8	9	10	11	12	13	14	15	16	17
8.0	83	75.0	9.0	33	25.0	32.0	9	12.5	-28.0	0.512
-1.9	31	43.3	-28.4	23	14.4	59.6	9	7.2	25.0	0.365
24.3	280	214.0	30.8	145	71.4	103.0	20	35.7	44.0	0.527
2.5	108	100.5	7.5	36	33.4	7.8	9	16.7	-46.1	0.525
-4.6	117	174.0	-32.8	84	58.0	44.8	25	29.0	-13.8	0.355
3.9	237	242.5	-2.3	100	80.7	23.9	48	40.4	18.9	0.470
-12.6	92	119.5	-23.0	36	39.8	-9.6	8	19.9	-59.9	0.440
2.2	153	138.0	10.2	48	46.0	4.4	18	23.0	-21.8	0.539
-17.8	104	135.0	-23.0	47	45.0	4.5	14	22.5	-37.8	0.469
27.4	76	57.8	37.4	43	19.6	119.6	5	9.8	-38.8	0.516
2.6	35	29.3	19.9	10	9.8	2.0	3	4.9	-32.2	0.584
26.2	204	158.0	28.4	90	52.6	66.1	28	26.3	1.9	0.514
15.7	206	162.5	27.4	70	54.0	29.6	27	27.0	0.0	0.553
13.9	144	104.0	38.5	29	34.0	-14.8	12	17.4	-30.8	0.608
8.2	74	67.5	9.6	24	22.5	6.7	11	11.2	-1.8	0.506
15.3	212	137.0	54.8	33	45.6	-27.7	18	22.8	-21.1	0.671
14.0	84	75.0	12.0	29	25.0	16.0	19	12.5	52.0	0.492
-14.1	48	60.0	-20.0	19	20.0	-5.0	3	10.0	-70.0	0.466
-0.6	76	80.0	-5.0	25	26.5	-5.7	13	13.3	-2.3	0.478

Literatur

- BORODA, M.G., K voprosu o metroritmičeski elementarnoj edinice v muzyke (Zur Frage der metrisch-rhythmischen Elementareinheit in der Musik). In: Soobščeniya AN GSSR 71, 1973, 745-748
- BORODA, M.G., O melodičeskoj elementarnoj edinice (Über die melodische Elementareinheit). In: "MAAFAT-75". Materialy pervogo vsesojuznogo seminaru po mašinnyj aspektam algoritmičeskogo formalizovannogo analiza muzykal'nyh tekstov (Materialien zum Ersten Allunions-Seminar zu Aspekten maschineller Verfahren bei der algorithmisch formalisierten Analyse musikalischer Texte). Erevan, Izdatel'stvo AN Armjanskoj SSR, 1977, 112-120
- BORODA, M.G., NADAREJŠVILI, I.S., ORLOV, Ju.K., ČITAŠVILI, R.Ja., O karaktere raspredelenija informacionnyh edinic maloju častoty v chudožestvennyh tekstach (Über den Charakter der Distribution von informationstragenden Einheiten mit geringer Häufigkeit in künstlerischen Texten). In: Semiotika i informatika 9, 1976, 23-34
- DETLOVS, V.K., O statističeskom analize muzyki (Über die statistische Musikanalyse). In: Latvijskij matematičeskij ežegodnik, Riga 1968, Nr. 3, 101-120
- FRUMKINA, R.M., Statističeskie metody izučeniya leksiki (Statistische Methoden der Untersuchung des Lexikons). Moskva, Nauka 1964
- FUCKS, W., Po vsem pravilam iskusstva (Nach allen Regeln der Kunst). In: BIRJUKOV, B.V., ZAZHIPOV, R.H., PLOTNIKOV, S.N. (Hrsg.), Iskusstvo i EVM (Kunst und Computer). Moskva, Mir 1975, 303-309 (vgl. FUCKS, W., Nach allen Regeln der Kunst. Stuttgart 1968 (Anm. d. Übers.))
- KAC, B., O nekotoryh čertach struktury variacionnogo cikla (Über einige Strukturmerkmale des Variationszyklus). In: Voprosy teorii i estetiki muzyki (Fragen der Musiktheorie und -ästhe-

tik). Vyp. 11, Leningrad 1972, 167

- MAZEL', L.A. CUKKERMANN, V.A., Analiz muzykal'nyh proizvedenij (Die Analyse musikalischer Werke). Moskva, Muzyka 1967, 393-490
- ORLOV, Ju.K., Častotnye struktury konečnyh soobščenij v nekotorych estestvennyh informacionnyh sistemach (Häufigkeitsstrukturen endlicher Mitteilungen in einigen natürlichen Informationssystemen). Kand. diss. Tbilisi 1975 (typescript), 59-87
- ORLOV, Ju.K., Statističeskaja struktura soobščenij, optimal'nyh dlja čelovečeskogo vosprijatija (Die statistische Struktur von Mitteilungen mit für die menschliche Wahrnehmung optimaler Ausformung). In: Naučno-tehničeskaja informacija. Serija 2, Nr. 8, 1970, 11-16
- ROJTERŠTEJN, M.I., Graf i matrica kak instrumenty ladovogo analiza (Graph und Matrix als Mittel zur Tonartanalyse). Muzykal'noe iskusstvo i nauka (Musikkunst und Wissenschaft). Vyp. 2, Moskva 1973, 175-189
- TJULIN, Ju.N. (Hrsg.), Muzykal'naja forma (Die musikalische Form). Moskva, Muzyka 1974, 46-49

DISKUSSION ZU M.G. BORODA,
HÄUFIGKEITSSTRUKTUREN MUSIKALISCHER TEXTE

G. Altmann

1. Es werden in der Arbeit mehrere Inferenzen über Datenverläufe gemacht, jedoch fehlen vorläufig objektivierte Testprozeduren. Ein einfacher Anpassungstest kann nachträglich durchgeführt werden. Es ist (trotzdem) zu erwarten, daß die Arbeit nicht nur Linguisten und Musikologen, sondern auch Statistiker stimulieren wird.

2. Es ist nicht evident, daß die Wiederholungsstruktur von F-Motiven mit der relativen Häufigkeit des häufigsten F-Motivs ($p_{\max}$) bzw. mit der Textlänge ursächlich im Zusammenhang steht. Wäre das der Fall, so müßten diese Größen bereits in der Ableitung des Zipf-Mandelbrotschen Gesetzes angesetzt werden. Es scheint, als ob die Parameter der Kurven mit Hilfe dieser Größen lediglich empirisch geschätzt wären, auch wenn diese Schätzungen linguistisch-musikalisch a posteriori sehr plausibel erscheinen. Andere Schätzer sind aber auch möglich. Daher muß die Aussage "Demnach ist die Struktur der Wiederholungen von Wörtern im literarischen Text abhängig von der Länge dieses Texts und von der Frequenz des in ihm am häufigsten vorkommenden Wortes" cum grano salis betrachtet werden.

3. Um von einem Gesetz (im Unterschied zu einer empirischen Generalisierung) sprechen zu können, muß folgendes geleistet werden:

a) Die Kurve muß aus einer Theorie abgeleitet werden (theoretische Validierung), wobei die Beziehung der Koeffizienten zu irgendwelchen Texteigenschaften a priori begründet werden soll. Die Theorie braucht nicht ausgereift zu sein, es reichen plausible Annahmen, die vorläufig die Rolle der Axiome übernehmen. Es wäre ratsam, Arbeiten russischer Autoren zu diesem Thema zugänglich zu machen.

b) Die Anpassungen müssen empirisch streng überprüft werden (empirische Validierung).

c) Die Anomalien, die bei Teiltexen entstehen, müssen später durch eine zusätzliche Annahme erfaßt werden. Die Zipf-Mandelbrotsche Kurve wird wahrscheinlich mehrere Verfeinerungen erfahren und mit der Zeit immer komplizierter werden. Ad-hoc-Hypothesen soll man jedoch nicht pflegen.

4. Die untersuchte Texterscheinung wurde unter eine Gesetzeshypothese subsumiert, was für eine wissenschaftliche Erklärung ausreichend wäre, wenn uns noch die musikologische Interpretation des Entstehungsmechanismus dieser Erscheinung bekannt wäre. Der Kern der künftigen theoretischen Arbeit, bei der Mathematiker, Musikologen und Psychologen zusammenarbeiten müssen, liegt genau in diesem Punkt. Die Mandelbrotsche Begründung des Zipfschen Gesetzes für rein linguistische Zwecke reicht nicht aus; die Annahmen, die zu ihm führen, sollten einen musikologischen (psychologischen) Hintergrund haben.

5. Die Tatsache, daß bei der Textbildung nicht nur "grammatische Regeln", sondern auch Makrogesetze wirken, die sich mathematisch erfassen lassen, ist die beste Voraussetzung für eine künftige Texttheorie, die sich nicht nur auf sprachliche Texte beschränkt. Die von Boroda gefundene Analogie hat einen unschätzbaren Wert für mindestens drei wissenschaftliche Disziplinen: Linguistik, Musikologie, Psychologie.

E. Fischer

Der vorliegende Aufsatz von Boroda vermittelt eine Reihe aufschlußreicher Ergebnisse und eröffnet durchaus Möglichkeiten, der Betrachtung musikalisch-kompositorischer Zusammenhänge neue Aspekte zu erschließen. Gleichwohl wird die Valenz der Resultate bisweilen dadurch gemindert, daß 'traditionelle' Beschreibungs- bzw. Kategorisierungsansätze bei der Auswahl und Interpretation der Exempla kaum Berücksichtigung finden. Zum Beispiel gründet die Behauptung, 'die Erfüllung des Zipf-Mandelbrotschen Gesetzes stehe in unmittelbarem und wesentlichem Zusammenhang mit dem musikalischen Werk als Ganzem' (S. 52), einzig auf Analysen von Kompositionen (S. 48-51: Präludien und Fugen, Sonaten), bei denen das Verhältnis zwischen vollständigem Text und Textaus-

schnitt keineswegs a priori mit demjenigen übereinzustimmen braucht, das in anderen 'zyklischen' Werken - wie Suiten oder Variationen-Folgen - vorherrscht. Auch wäre die statistische Argumentation, nach der bei Texten mit annähernd gleicher Länge und annähernd gleichen Werten für $p_{\max}$ die Häufigkeitsstrukturen kongruieren (S. 53f.), überzeugender gelungen, wenn dem Graphen-Vergleich in Abb. 6.1 statt eines zweiten Rondos mindestens eines der anderen untersuchten Stücke gedient hätte, an denen doch ähnliche Verhältnisse zu beobachten sein sollen (S. 55): nur derart wäre der plausible Einwand zu entkräften gewesen, daß die vorgeführten Korrelationen - unbeschadet der 'grundlegenden stilistischen Unterschiedlichkeit der Texte' (S. 54) - weniger auf der 'annähernd gleichen Länge' der Kompositionen, sondern vielmehr auf der transepochalen Konstanz eines generisch-typologischen Modells beruhen könnten. (Für den zweiten, komplementären Untersuchungsschritt (vgl. Abb. 6.2) wählt der Autor bezeichnenderweise Werke aus, die sehr verschiedenen Genera zugehören.)

Gravierender als die genannten Kritikpunkte sind jene Bedenken, die das Segmentierungsverfahren - soweit es in diesem Aufsatz eingeführt und erläutert ist - hervorruft. Denn eine 'musikalische Elementareinheit', die sich 'im Ausnahmefall mit einem Motiv oder Submotiv deckt' (S. 44), ist für den Musikwissenschaftler ebenso problematisch, wie es für den Linguisten eine sprachliche Elementareinheit wäre, die eher zufälligerweise mit Silben oder Wörtern korrespondierte. Einesteils zwingt der Sachverhalt, daß 'Motiv' oder 'Submotiv' häufig äquivok definiert werden (S. 42), nicht unbedingt dazu, auf eine musikwissenschaftlich befriedigende Bestimmung dieser Einheiten zu verzichten; andernfalls müssen Einheiten mit fixierter Länge keinesfalls stringent zu einer 'unnatürlichen' (S. 41) Aufgliederung musikalischer Texte führen: gerade das geregelte Akzentsystem taktgebundener Musik - in dieser Perspektive hätte sich ein fruchtbarer Vergleich mit der metrischen Konstitution poetischer Systeme angeboten - sollte es wohl erlauben, variable, dem künstlerischen Medium adäquate Zeiteinheiten auszugrenzen, die zueinander in sinnvollen Proportionen stehen und nicht - wie die doch auch primär nach metrisch-rhythmischen Ordnungsprinzipien selektierten

'F-Motive' - ganze Takte und gleichermaßen isolierte Einzeltöne umfassen dürften. Borodas Gliederung des Čajkovskij-Beispiels (Abb. 3.d) zeigt, daß hier höchst relevante Rekurrenzen kompositorischer Mikrostrukturen vernachlässigt werden, während nicht nur im Šostakovič-Beispiel (3.f) Klänge als identische 'F-Motive' zu klassifizieren sind, obschon ihnen keine musikalische Gestaltqualität zukommt. Deshalb erscheint letztlich die Vermutung nicht unbegründet, daß durch die vorgenommene Textsegmentierung eine weitere, dem Verfahren selbst inhärierende 'Kompensationserscheinung' begünstigt werden könnte: In Werken intentional angelegte Wiederholungsstrukturen mögen unberücksichtigt bleiben, solange dafür zusätzliche Entsprechungen mit Hilfe der 'F-Motive' konstruiert werden.

W.-D. Schäfer

Ungeachtet der Schwierigkeiten, die tatsächliche Relevanz der musikalischen Bedeutung der F-Motive zu begründen, sind die Ergebnisse von Boroda beeindruckend. Nicht so sehr die trivial anmutende These, daß es eine geringe Anzahl häufig wiederholter F-Motive und eine beträchtliche Anzahl seltener F-Motive gibt, verblüfft, als vielmehr die gute Beschreibbarkeit mittels eines mathematisch formulierten Gesetzes. Einerseits möchte ich hier auf die Übertragbarkeit des verallgemeinerten Zipf-Mandelbrotschen Gesetzes auf andere musikalische Parameter eingehen, andererseits die hier vorgestellten Ergebnisse aufgrund meiner Forschungsergebnisse kommentieren.

Bemerkenswert erscheint zunächst, daß die festgestellten Abweichungen von den theoretisch ermittelten Werten nicht eine historische Entwicklung oder Gesetzmäßigkeit widerzuspiegeln scheinen, was Boroda auch veranlaßt, hier von einem allgemeinen Gesetz zu sprechen. Dies kann einmal darauf deuten, daß sich die ausgewählten Werke gerade in Hinsicht ihrer Motivverarbeitungstechnik kaum unterscheiden - somit wäre hier ein Charakteristikum für die Musik der Barockzeit bis zum ausgehenden 19. Jahrhundert jenseits epochaler Stilmerkmale gewonnen - , auf der anderen Seite besteht aber freilich auch die Gefahr, daß die Konstruktion des F-Motivs

selbst sich allzu sehr in den gewonnenen Ergebnissen widerspiegelt, da das F-Motiv als eine Fiktion definiert wird, die aus dem musikalischen Material der klassisch-romantischen Epoche abstrahiert wurde.

Eine Anwendung des Zipf-Mandelbrotschen Gesetzes in der Gestalt (4), S. 39, auf einen anderen Parameter der Musik, die Instrumentation, führt uns durchaus auf Ergebnisse, die eine historische Interpretation verlangen. Im Zuge meiner Forschungsarbeiten¹ wurden einige ausgewählte Orchestersätze der Klassik und der Spätromantik auf ihr Klangrepertoire hin untersucht. Das Äquivalent zu meinem Begriff 'Klangrepertoire' ist bei Boroda das 'Motivinventar'.

Eine Überprüfung hat ergeben, daß die Repertoireschätzung mittels der Formel (4) bei den klassischen Sinfoniesätzen ganz erheblich von den tatsächlichen Repertoiregrößen abweicht, bei den spätromantischen Werken dagegen ist eine sehr viel bessere Übereinstimmung zu konstatieren.

Tabelle 1²⁾:

Komponist	n	n*	δ
Mozart	72	270	-73.3%
Haydn	90	219	-58.9%
Beethoven	154	433	-64.4%
Wagner	186	232	-19.8%
Schönberg	180	181	- 0.6%
Mahler	656	651	- 0.8%

Dieses Defizit an Repertoireelementen innerhalb des klassischen Sinfoniesatzes läßt eine historische Deutung zu: als selbständiger Parameter tritt die Instrumentation in ihren vielfältigen Möglichkeiten erst im Laufe der Romantik in das Bewußtsein und wird voll in der Spätromantik genutzt.

In diesem Zusammenhang stellt sich allerdings auch die Frage, inwieweit Abweichungen von den erwarteten Repertoiregrößen noch tolerabel sind. Leider finden sich bei Boroda in dieser Hinsicht keine Erklärungen. Ein Blick auf die bei Boroda abgedruckte Tabelle läßt sofort erkennen, daß die Ergebnisse einem Prüfverfahren mit der Schärfe des χ^2 -Tests sicher nicht gewachsen sind.

Man wird in diesem Bereich der musik-literaturwissenschaftlichen Untersuchungen zu weiteren Formulierungen kommen müssen. Der rein anschauliche Vergleich von Kurven mag zwar vieles plausibel machen, ersetzt aber eine mathematische Begründung nicht, hinterläßt er doch immer beim Leser ein unbefriedigendes Gefühl.

Andere wichtige Erkenntnisse lassen sich im übrigen aus meinen Arbeiten bestätigen:

1. die Übertragung der Gesetzmäßigkeiten vollständiger Werke auf Teilbereiche, Ausschnitte, sogar auf eine Menge zufällig ausgewählter Klangereignisse der Komposition ist schwierig, wenn nicht gar unmöglich. Das führt zu dem großen Problem, daß bisher keine geeigneten Schätzmethoden auf Größe von Repertoires oder von Verteilungen gefunden werden konnten. Das wiederum bewirkt, daß derartige Untersuchungen immer sehr aufwendig sind, da die Kompositionen vollständig ausgezählt werden müssen;
2. die besondere Bedeutung der doppelten Okkurrenz eines Ereignisses in musikalischen Werken läßt sich auch hier vielfach belegen (zur Begründung vgl. den Beitrag von Wildgruber). Von den 656 Klangereignissen des Repertoires bei Mahler treten 188 nur einmal auf, dagegen 206 zweimal. Zusammen genommen erst erreichen sie annähernd den theoretischen Wert für n_1 . Dies zeigt deutlich, wie auch die Ergebnisse von Boroda, daß die Schätzungen für n_m noch einer erheblichen Verbesserung bedürfen.
3. Eine Übereinstimmung ergibt sich ebenfalls mit der Beobachtung, daß in längeren musikalischen Texten das Repertoire schneller anwächst, wie ein Vergleich der Kurven von Schönberg und Mahler in der Abbildung 1 deutlich zeigt. Bei annähernd gleicher Länge (Beispiele von Schönberg und Wagner) scheinen sich allerdings auch stilistische Unterschiede bemerkbar zu machen.

Interessanter als diese Bestätigung scheint aber die relativ gleiche Verteilung der Repertoirezunahme über die gesamte Komposition zu sein im Vergleich zu Mahler, bei dem das Anwachsen des Repertoires sehr unregelmäßig verläuft. Zum besseren Vergleich ist das Repertoire in prozentualen Angaben in der Abbildung 2 eingetragen.

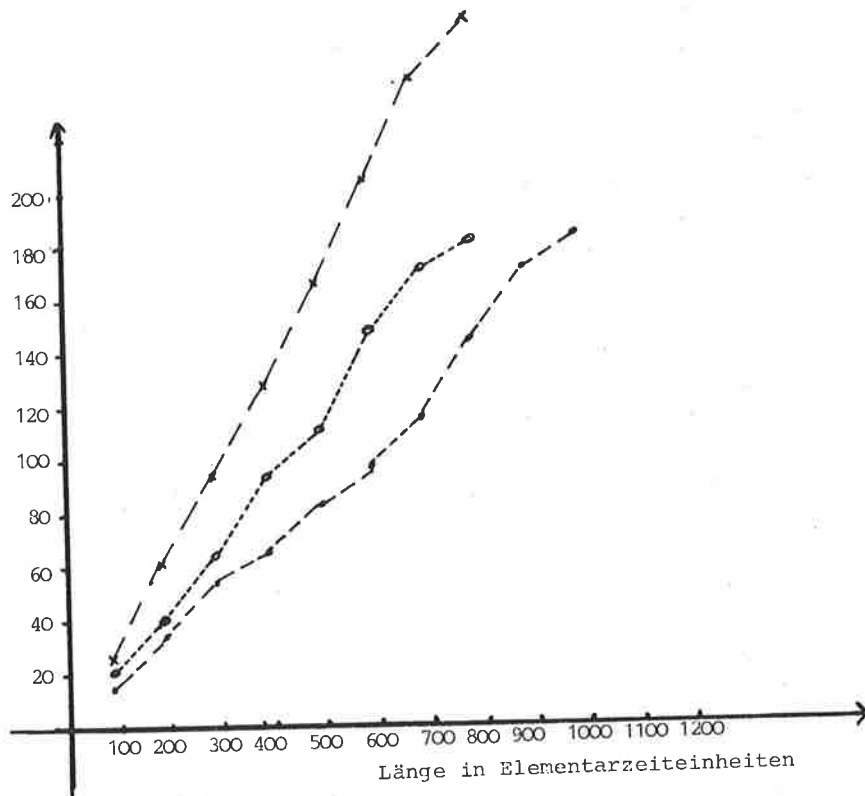


Abb.: 1

- Wagner 940 Einheiten
- Schönberg 768 Einheiten
- x Mahler 3322 Einheiten

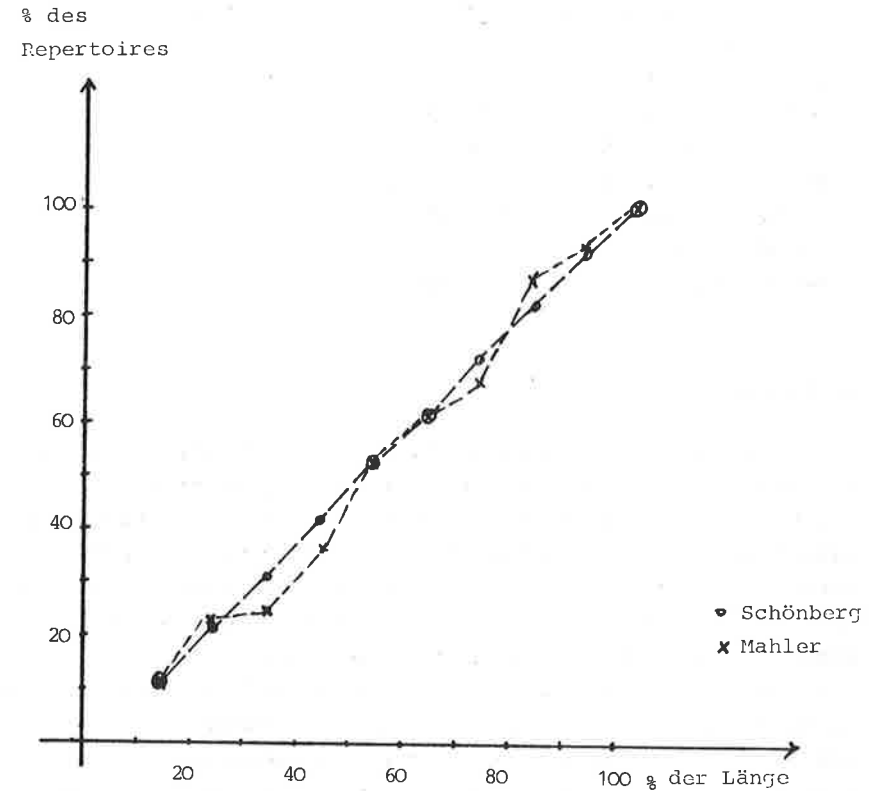


Abb.: 2

Anmerkungen

- 1) Die Ergebnisse werden demnächst in meiner Dissertation über 'Möglichkeiten quantitativer Instrumentationsanalyse' veröffentlicht.
- 2) Da die Länge der Klangereignisse verschieden interpretiert werden kann, einerseits auf eine absolute Elementarzeiteinheit bezogen, andererseits bis zum nächsten Besetzungswechsel - führen verschiedene Rechenverfahren auf unterschiedliche Ergebnisse. Die Tendenz jedoch, daß in der Klassik zunächst viel zu kleine Repertoires auftreten, die dann in der Spätromantik sich aber der Größen-

ordnung des Zipf-Mandelbort'schen Gesetzes annähern, ist jeweils die gleiche.

Im einzelnen handelt es sich um die Werke:

Mozart, Hafner-Sinfonie, 1. Satz

Haydn, Londoner Sinfonie Nr. 104, 1. Satz

Beethoven, Sinfonie Nr. 1, 1. Satz

Wagner, Vorspiel zu 'Parsifal'

Schönberg op. 16, 1

Mahler, Sinfonie Nr. 5, 1. Satz

J. Wildgruber

Die Übertragbarkeit des Zipf-Mandelbrotschen Gesetzes auf musikalische Texte vom Barock bis zur gemäßigten Moderne ist ohne Zweifel eine bemerkenswerte Entdeckung; stiftet sie doch nicht nur einen Zusammenhang innerhalb der Epoche bürgerlicher Musikentfaltung, sondern auch zwischen literarischen und musikalischen Texten. Umso interessanter dürften jene Abweichungen von dieser Gesetzmäßigkeit sein, die im Medium Musik begründet sind.

Mit Recht beschränkt sich der Verfasser auf einen Abschnitt der Musikgeschichte, in dem sich die Musik zunehmend von der Wortgebundenheit entfernt, sich vom Vokalstil emanzipiert, und Termini wie Takt, Rhythmus, Motiv, Thema sich als adäquate Beschreibungsmerkmale erwiesen haben; gerade für diesen Bereich aber hätte sich das 'Prinzip des zweimaligen Vorkommens eines F-Motivs' erwarten lassen. Denn im Gegensatz zum literarischen Text, in dem in aller Regel jede Buchstabenfolge ein wohldefiniertes Zeichen darstellt, muß sich die Musik ihre Zeichen erst schaffen. Die einfachste Möglichkeit, einer beliebigen Tonfolge (oder einem beliebigen F-Motiv) Zeichencharakter zu verleihen, ist offensichtlich ihre Bestätigung durch Wiederholung. In dieser Form prägt sie den Typus der Themen mit satzartiger Entwicklung ebenso wie jenen, der sich als Periode darstellt, d.h., sie ist ein Stilmerkmal von der Vorklassik bis zum Neoklassizismus. Bei einem hypothetisch-deduktiven Vorgehen hätte sich für diesen zunächst sehr grob umrissenen Bereich musikalischer Texte eine Korrektur des Zipf-Mandelbrotschen Gesetzes von vornherein als notwendig

erwiesen. In einem solchen Kontext erschiene es eher als über-raschend, wenn F-Motive nur einmal auftreten: es könnte sich ebenso gut um musikalischen Leerlauf handeln, um beliebig austauschbare Floskeln, wie auch um Bestätigungen des Zeichencharakters, die nur im Rahmen der Identität von F-Motiven nicht erkennbar sind. Als Beispiel für derartig korrespondierende bzw. antithetische Motivwiederholungen seien die folgenden Takte genannt:

Siegmund.

sm. Friedmund darf ich nicht heißen; Frohwalt möchte wohl sein: doch Weh-walt

sm. muß ich mich nennen.

Nicht zufällig ist hier ein Wagner-Beispiel gewählt worden: Eine solche erfahrbare, im Rahmen der F-Motiv-Konstruktion jedoch nicht meßbare Identität der motivischen Gestalt verweist auf Mahler und über diesen hinaus auf die Reihentechnik, in der die Frage nach einer 'Identität' von Motiven überaus problematisch wird.

M.G. Boroda

Zu G. Altmann

(a) Wie Altmann bemerkt, fehlen in meiner Arbeit statistische Prozeduren zur Überprüfung der (statistischen) Hypothese, daß in musikalischen Texten das Zipf-Mandelbrotsche Gesetz gilt. Ich bin einverstanden, daß eine solche Überprüfung sehr wünschenswert ist. Man kommt auch leicht auf die Idee, irgendein bekanntes statistisches Kriterium zu verwenden, z.B. den Chi-Quadrat-Anpassungstest. Aber dieser Gedanke stößt an einen ersten Einwand. Das Problem liegt darin, daß die bekannten statistischen Kriterien zur Beurteilung der Übereinstimmung einer empirischen mit einer theoretischen Häufigkeitsverteilung auf der Annahme basieren: "Je größer die Stichprobe aus einer Grundgesamtheit, desto wahr-

scheinlicher ist es, daß die in der Stichprobe beobachteten Häufigkeiten von Elementen ihren wahren Wahrscheinlichkeiten in der Grundgesamtheit hinreichend ähnlich werden". Oder, etwas gröber gesagt, mit dem Anwachsen der Stichprobe nähern sich die Häufigkeiten ihrer Elemente ihren Wahrscheinlichkeiten.

Aus Platzgründen ist dies nicht der geeignete Ort, um das Problem zu erörtern, warum man bei musikalischen (und auch literarischen) Texten schwerlich von Wahrscheinlichkeiten der Elemente sprechen kann, so daß es in diesen Fällen äußerst schwierig ist, die Grundgesamtheit zu bestimmen (in Bezug auf literarische Texte vgl. die Arbeiten von Orlov 1976, 1977 u.a.). Dazu ist noch zu bemerken, daß in den von uns untersuchten musikalischen Texten und Abschnitten aus ihnen keine bessere Übereinstimmung mit dem Zipf-Mandelbrotschen Gesetz in Stichproben größeren Umfangs beobachtet wurde. Wie aus Abb. 5.1 und 5.2 ersichtlich, hat ein kleinerer vollständiger Text (vgl. die Häufigkeitsgraphik von Bachs Präludium und Fuge Nr. 2 aus dem "Wohltemperierten Klavier" Bd. 2 und die des Abschnitts (Contrapunctus Nr. 8) aus seiner "Kunst der Fuge") eine bessere Übereinstimmung mit dem Gesetz gezeigt als ein längerer Abschnitt aus einem Text. Der vollständige Text erfüllt das Zipf-Mandelbrotsche Gesetz zufriedenstellend, wie man nach der Häufigkeitsgraphik urteilen kann, während der Abschnitt derselben Länge eine offensichtliche Abweichung der beobachteten Häufigkeiten der F-Motive von der theoretischen Kurve aufweist. Dieselben Erscheinungen hat man auch in literarischen Texten beobachtet (dort wurden außer Textabschnitten auch Zusammenstellungen mehrerer Texte zu einer Stichprobe untersucht). Auch diese Zusammenstellungen zeigten signifikante Abweichungen der Worthäufigkeitsstruktur von dem Zipf-Mandelbrotschen Gesetz in seiner Form (3) (vgl. zu diesem Problem auch Arapov, Efimova, Šrejder 1975).

Aus diesem Grund erweisen sich die üblichen statistischen Testprozeduren zur Überprüfung des Zusammenhangs zwischen der Organisation der Wiederholungen von F-Motiven in musikalischen Texten (oder von Wörtern in literarischen Texten) und dem Textganzen als ungeeignet. Für Daten dieser Art müßten spezielle Methoden ausgearbeitet werden, die die Endlichkeit und die Abgeschlossenheit

des Textes berücksichtigen. Soweit uns bekannt ist, hat man gerade angefangen solche Methoden zu entwickeln und die Forschung ist noch bei weitem nicht abgeschlossen.

(b) Wie Altmann bemerkt, ist es nicht evident, daß die Struktur der Wiederholung von F-Motiven ursächlich mit der Textlänge und mit der relativen Häufigkeit $p_{\max}$ des häufigsten F-Motivs zusammenhängt. Mit anderen Worten, Altmann vertritt die Meinung, daß es möglicherweise andere, "wesentlichere" Parameter der Organisation der Wiederholung von F-Motiven geben kann.

Dieser Gedanke ist interessant. Dazu jedoch zwei Bemerkungen. Erstens, das mathematische Modell eines Prozesses ist desto besser, je weniger Anpassungsparameter es (ceteris paribus) hat. In dieser Hinsicht hat die aus Orlovs Modell resultierende Form des Zipf-Mandelbrotschen Gesetzes eine wertvolle Eigenschaft: sie hat in der Tat keine Anpassungsparameter, da K , B und γ der kanonischen Formel von Mandelbrot in Orlovs Modell Funktionen zweier natürlicher und leicht berechenbarer Textcharakteristika sind, nämlich der Textlänge N_0 und der relativen Häufigkeit $p_{\max}$ des häufigsten Textelements. Die Tatsache, daß ein solches parameterloses Modell die Wiederholungsstruktur in abgeschlossenem literarischen oder musikalischen Text adäquat beschreibt, ist ein Beweis der Stärke des Modells. Zweitens, wie soll man die zusätzlichen oder "wesentlichen, ursächlichen" Parameter der Wiederholungsstruktur von F-Motiven suchen? Es ist recht hoffnungslos und äußerst arbeitsaufwendig die Methode von "trial and error" zu verwenden. Es ist schwer, bei dieser Suche die Daten der musikwissenschaftlichen Untersuchungen zu benutzen, da ihre Resultate in der Regel keinen allgemeinen Charakter haben und gewöhnlich in einer verschwommenen, rein "qualitativen" Form ausgedrückt werden, so daß ein Übergang zu strukturellen Charakteristika extrem kompliziert wäre. Abgesehen von diesen Umständen ist Altmanns Idee über die Existenz "ursächlicher" Parameter der Wiederholungsstruktur der F-Motive wertvoll. Bei weiteren Untersuchungen der Wiederholungsstruktur in musikalischen Werken sollte man diesen Gedanken unbedingt in Betracht ziehen.

(c) Zu Altmanns Punkten (3) und (4) und seinen Empfehlungen folgendes:

(c1) Die Häufigkeitskurve, mit Hilfe derer ich die Häufigkeitsstruktur der F-Motive im musikalischen Text analysiere, wurde in der Tat aus theoretischen Annahmen abgeleitet (vgl. Orlov 1975 u. a.). Die Koeffizienten dieser Kurve - die Textlänge und $P_{\max}$ - spielen auch bei der Rezeption des Textes eine wichtige Rolle. Die Suche nach den Prinzipien der Texterzeugung, die die Realisierung des Zipf-Mandelbrotschen Gesetzes sichern, stellt eine äußerst wichtige Aufgabe dar. Sie zog schon früher die Aufmerksamkeit der Forscher auf sich - so hat B. Mandelbrot seine Formel aufgrund der Theorie der optimalen Kodierung abgeleitet -, gegenwärtig verstärkt sie sich, sowohl im Hinblick auf mathematische Modelle als auch im Hinblick auf die psychologische Interpretation des Zipf-Mandelbrotschen Gesetzes. Die Besprechung dieser interessanten Probleme hätte den Rahmen meiner Arbeit gesprengt. Ich verweise auf die Arbeiten von Mandelbrot (1957), Arapov, Šrejder (1977), Orlov (1976, 1978), Price (1978), Boroda, Orlov (1978).

(c2) Bezüglich der musikologischen Interpretation meiner Resultate ist die Arbeit in dieser Richtung äußerst interessant und wurde bereits gestartet. Vgl. darüber meine Dissertation (Boroda 1979). Das letzte Kapitel ist speziell den Fragen der Interpretation der entdeckten quantitativen Gesetzmäßigkeiten der F-Motiv-Wiederholungen aus der musikologischen Sicht gewidmet. Dort wurde gezeigt, daß zwischen den vom musikologischen Standpunkt aus wichtigen Eigenschaften eines musikalischen Werks (Ausgeglichenheit der musikalischen Form, starke gegenseitige Abhängigkeit zwischen der Makro- und der Mikroebene des Werkes, strukturelle Einheit des Zyklus usw.), den Gesetzmäßigkeiten, die auf der Ebene der F-Motive beobachtet werden können, und speziell der Realisierung des Zipf-Mandelbrotschen Gesetzes im musikalischen Text, ein wesentlicher Zusammenhang besteht. An derselben Stelle wird in Form einer Hypothese die Rolle des Zipf-Mandelbrotschen Gesetzes auf der Ebene der F-Motive bei der Entwicklung musikalischer Formen in der europäischen Musik der letzten drei Jahrhunderte besprochen. Aus Umfangsgründen konnten leider die Fragen der musikologischen Interpretation in den Aufsatz in diesem Band nicht aufgenommen werden.

Zu E. Fischer

(a) Wie Fischer erwähnt, werden in meinem Aufsatz solche musikalischen Formen (Sonate, Symphonie, Präludium und Fuge) untersucht, in denen sich ein Abschnitt wie ein Abschnitt verhalten muß, während es in Suiten und einigen anderen zyklischen Formen nicht unbedingt so sein muß. D.h. Fischer meint offensichtlich, daß ein Abschnitt aus den von mir untersuchten Zyklen dem Zipf-Mandelbrotschen Gesetz nicht folgen dürfte und daher die von mir gezeigte Tatsache, daß das Gesetz im vollständigen musikalischen Text gilt und in einem Abschnitt nicht gilt, keine Überraschung sei. Obwohl Fischer diesen Gedanken nicht explizit äußert, ist seine Kritik offensichtlich gegen diesen Punkt gerichtet.

Die Kritik scheint mir aus folgendem Grund unbegründet. Welches ist denn in der Musiktheorie eines der wichtigsten Merkmale für die Abgeschlossenheit eines musikalischen Werkes (hier meinen wir nicht die Charakteristika wie "dramaturgische Ganzheit", "Vollständigkeit der figuralen Entfaltung" usw., sondern strukturelle Eigenschaften und Thematismus)? Die thematische Selbständigkeit und die Abgeschlossenheit der Form (unter der letztgenannten Eigenschaft versteht man gewöhnlich die Realisierung einer der geschlossenen Formen im gegebenen Werk: der Periode, der dreigliedrigen Form, der Sonatenform usw.). In der Regel ist ja - und in unseren Untersuchungen war es immer so - ein Teil der Sonate oder Symphonie, Fuge und Präludium im polyphonischen Zyklus sowohl thematisch selbständig als auch abgeschlossen in der Form! Außerdem sind diese Teile hinreichend entfaltet und weisen eine intensive musikalische Entwicklung auf. Um mich etwas mehr "qualitativ" auszudrücken, jeder Teil des Zyklus - des sonaten-symphonischen, des Präludium-Fuge-Zyklus u.a. - hat seinen eigenen figural-emotionalen Charakter und verwirklicht seine Stimmung (denken wir an den ziemlich stark reglementierten klassischen symphonischen Zyklus)! Und trotzdem spüren wir intuitiv, daß ein Teil des Zyklus eben nur ein Teil ist, etwas in bestimmter Hinsicht noch nicht Abgeschlossenes. Soweit mir bekannt, hat eben die Analyse auf der Ebene der F-Motiv-Wiederholungen zum ersten Mal erlaubt, diese Tatsache quantitativ von einem rein strukturellen Standpunkt aus

zu zeigen. Ein Teil des Zyklus einer Kompositionsform, die makrostrukturell gesehen ganz selbständig ist, erscheint vom Standpunkt der F-Motiv-Struktur als unbeeendet, als eingeordnet in den "Wiederholungsplan" des abgeschlossenen Werkes. Daher kann ich mit Fischers Kritik in diesem breiten Sinn nicht übereinstimmen. Wenn ich Fischers Worte als den Wunsch auffasse, daß ich auch den Suitenzyklus durchzählen sollte (zuzugeben wären auch der Variationenzyklus und andere Formen), so stimme ich damit voll überein. Solche Untersuchungen wären sowohl interessant als auch notwendig.

(b) Wie Fischer bemerkt, sollte neben der Häufigkeitsgraphik eines Rondos auf der Abb. 1 (wo zwei Rondos, Beethovens und Kabalevskijs, dargestellt sind) noch eine andere Form gezeigt werden, sonst entsteht der Eindruck, daß die Nähe der Häufigkeitskurven einfach durch die Invarianz (durch den transepochnalen Charakter) der Rondo-Form hervorgerufen wird.

Dieser Einwand wäre berechtigt, wenn es nicht den Umstand gäbe, daß die Rondo-Form hinsichtlich ihres Konstruktionsschemas nicht so "transepochnal" ist, wie sie scheint. Im Gegenteil, die Evolution dieser Form in Bezug auf die Entfaltung der Episoden, auf den Grad ihrer Selbständigkeit und Kontrastierung im Verhältnis zum Refrain usw., war so bedeutend, daß man bei der Analyse der musikalischen Formen von historisch unterschiedlichen Typen des Rondos spricht (vgl. z.B. Mazel' 1978). Das Rondo von Beethoven (op. 51, Nr. 1) und das Klavierrondo von Kabalevskij (op. 59) sind in ihrer Form überaus unähnlich. Ihr Kompositionsplan, vom Standpunkt der musikalischen Formenlehre, ist wie folgt (mit großen Buchstaben werden die großen Formteile bezeichnet; unterschiedliche Buchstaben bezeichnen thematisch unterschiedliches Material):

Kabalevskij:

Einleitung, A, B (= A_{ver}, Verarbeitung mit Motiven aus A), A' (Reprise-Refrain, verkürztes und variiertes A), C (= α, β, α, β', α' - "inneres Rondo"), B' (Scheinreprise mit Verarbeitungselementen; danach Motive aus A_{ver}), A'', Coda (Material aus C und A).

Schema: Einleitung, ABA'C'B'A'', Coda, oder: Einleitung, AA_{ver}A'
BA_{ver}A'' ≅ ABA

Beethoven:

A, B (neues Material), A' (verkürztes A), C (= αβα'), AC (Scheinreprise, danach Elemente aus C; Durchführungsarbeit), A'' (verkürztes A), Coda (mit den Verarbeitungselementen).

Schema: ABA'CA'', Coda ≠ ABA

Weiter, wenn ich von einem Unterschied der Häufigkeitsstrukturen in unterschiedlich langen Texten (mit ähnlichem p_{max}) spreche, dann führe ich nicht nur Werke mit unterschiedlichen Formen, sondern auch (vgl. Abb. 7) Werke mit gleichen Formen von demselben Komponisten auf: zwei Sonaten von Chopin. Der Vergleich bezeugt die Richtigkeit der These von der Abhängigkeit der Häufigkeitsstruktur der F-Motive von der Länge des musikalischen Textes. In diesem Sinne soll man die Abb. 6 und 7 und die aus ihnen gezogenen Schlußfolgerungen als Glieder einer Kette betrachten. Dies hat Fischer, wie mir scheint, nicht voll berücksichtigt.

(c) Fischers wesentliche Kritik richtet sich schließlich gegen die F-Motive. Seine Bemerkungen wurden anscheinend dadurch hervorgerufen, daß in der vorliegenden Arbeit das Problem der elementaren musikalischen Einheit, das notgedrungen zur Bestimmung des F-Motivs führte, nicht ausführlich besprochen werden konnte. Ich beantworte Fischers Anmerkungen der Reihe nach.

(c1) Wie Fischer schreibt, muß dem Musikologen das F-Motiv problematisch erscheinen, da es einmal dem Motiv, einmal dem Teilmotiv entspricht. Diesen Einwand kann man aus dem Grunde nicht annehmen, weil das Motiv (und das Teilmotiv als sein Teil) in der Musikologie so uneindeutig definiert und behandelt wird, daß es nicht möglich ist, es als eine strukturelle Einheit zu betrachten wie z.B. die Silbe oder das Wort in der Linguistik. Folglich kann man sich bei der Bestimmung anderer Einheiten nicht an das Motiv oder das Teilmotiv halten, als ob sie irgendwelche "Maßeinheiten" wären. Man vergleiche z.B. die folgenden Definitionen:

(1) Das Motiv soll man als ein musikalisch-inhaltliches Ausdruckselement des Themas (oder des thematischen Materials überhaupt) und nicht als eine metrisch-strukturelle Einheit seines Aufbaus betrachten... Unter Motiven soll man Intonationswendungen verstehen, die in dem Maße, in dem sie dem Thema charakteristische

Züge verleihen, ihre eigene Ausdrucksbedeutung haben (Tjulin et al. 1974: 46-48).

(2) Das Motiv ist der kleinste Teil eines musikalischen Gedankens, der den Status einer inhaltlichen (ausdrucksvollen) und konstruktiven (Aufbau-) Einheit hat (Mazel' & Cukerman 1967: 552).

(3) Das Motiv ist der kleinste Teil der Melodie, der harmonischen Tonfolge, der inhaltlich geschlossen ist und unter anderen analogen Konstrukten erkannt werden kann. In der Regel enthält das Motiv einen starken Teil und gleicht daher oft einem Takt. Bei bestimmten Tempo-, Maß- und Fakturbedingungen des musikalischen Werkes sind auch längere Zweitaktmotive möglich (Musik- enzyklopädie 1976: 696).

Diese Definitionen - und man kann noch eine ganze Reihe analoger Definitionen aus anderen Quellen, darunter auch von westlichen Autoren (vgl. z.B. die Arbeiten von Riemann 1929: 215, Lobe 1887: 79, enzyklopädisches Wörterbuch von Grove 1927: 530, Seeger 1966: 119 u.a.) anführen - zeigen deutlich, daß das Motiv uneindeutig definiert wird, daß man aus diesen Definitionen keine klaren und objektiven operationalen Definitionen formulieren kann, die es erlauben würden, die Motive aus einer Melodie zu segmentieren, eine Melodie in Motive zu zerlegen. In einer Reihe von Arbeiten der letzten Jahre wird sogar die Möglichkeit, eine beliebige Melodie in Motive zu zerlegen bestritten (vgl. z.B. Tjulin et al. 1974). Der Status des Teilmotivs ist vollkommen analog. Es ist offensichtlich, daß man bei der quantitativen Analyse musikalischer Werke solche Einheiten nicht verwenden kann; es ist völlig illusorisch mit Hilfe von Definitionen der oben angegebenen Art formale Äquivalente des Motivs oder des Submotivs konstruieren zu wollen.

Gerade aus diesem Grund entstand die Notwendigkeit eine elementare melodische Einheit zu definieren, die man in der Melodie rein nach strukturellen Merkmalen absolut eindeutig identifizieren kann. Zu bemerken ist, daß das F-Motiv die Melodie nicht nur eindeutig, sondern auch lückenlos vom Anfang bis zum Ende zerlegt. Mit Motiven und Teilmotiven kann man eine solche Zerlegung nicht erreichen. Sie ist aber für die Analyse eines vollständigen Musiktextes unbedingt erforderlich.

Es ist offensichtlich, daß man das F-Motiv (oder eine andere streng definierte Einheit), strenggenommen, mit dem Motiv und Submotiv nicht vergleichen kann, da diese im Grunde nicht definiert sind und nicht eindeutig aufgefaßt werden. Trotzdem ist es inhaltlich berechtigt, die Zerlegung einer Melodie in F-Motive mit jenen konkreten Zerlegungen zu vergleichen, die man in den musikologischen Arbeiten findet.

Ich möchte die Neuartigkeit des F-Motivs nachdrücklich unterstreichen. Die Zerlegung der Melodie in F-Motive lehnt sich an ihre metrisch-rhythmische Organisation an, folgt dieser Organisation, den in ihr entstehenden Unterbrechungen, den metrischen Akzenten. Gerade deswegen kann sich ein F-Motiv über einen ganzen Takt erstrecken oder mit einem einzigen Ton zusammenfallen (nebenbei bemerkt: die Länge des F-Motivs ist in der Tat scharf begrenzt - es übersteigt nicht fünf Töne). Leider kann ich hier auf eine ausführliche Diskussion des Problems der Elementareinheit und der Eigenschaften des F-Motivs nicht eingehen - ich verweise lediglich auf meinen Aufsatz (Boroda 1977), dessen Veröffentlichung in dieser Schriftenreihe geplant wird¹⁾.

(c2) Bezüglich Fischers Bemerkung, daß Einheiten mit konstanter Länge nicht zu einer unnatürlichen Zerlegung eines musikalischen Textes führen.

Mit dieser Ansicht kann ich keineswegs übereinstimmen. Hier eines der zahlreichen Beispiele:

W.A. Mozart. Sonate für Violine und Klavier, e-moll (1. Teil)



Ähnliche Beispiele kann man in beliebiger Menge finden: immer, wenn die Zerlegung eine starke rhythmische Unterbrechung, die rhythmische Verbindung, schneidet, entsteht der Eindruck der Unnatürlichkeit (nebenbei bemerkt, bei der Zerlegung in F-Motive kann es derartige "Einschnitte" nicht geben, weil eine der Grund-

lagen des F-Motivs die rhythmische Anlehnung eines Tones an den folgenden längeren Ton ist)². Ausführlicher über Einheiten mit konstanter Länge vgl. meine Dissertation (Boroda 1979).

(c3) Nun zu Fischers Bemerkung, daß meine Gliederung der Melodie des Themas aus dem II. Teil von Čajkovskis 5. Symphonie die höchst relevante Rekurrenz der kompositorischen Mikrostrukturen gänzlich vernachlässigt.

Diese Kritik ist sehr kategorisch, jedoch kann man gegen sie schwer argumentieren, da es unklar ist, was Fischer unter der "Rekurrenz kompositorischer Mikrostrukturen" versteht. Deswegen zeige ich hier, daß gerade eine F-Motiv-Analyse dieses Themas von Čajkovskij solche Besonderheiten zu entdecken erlaubt, die man ohne Hilfe der F-Motive nicht beobachten kann.

Man kann diese Melodie sehr gut in drei größere Segmente zerlegen:



Man kann noch bemerken, daß das dritte Segment eine eigenartige Summationsfunktion erfüllt. Das ist leider alles.

Betrachten wir nun die F-Motiv-Zerlegung. Die ersten zwei F-Motive decken sich mit den Segmenten A und A'. Weiter kann man aber eine interessante Erscheinung beobachten: Das dritte F-Motiv gibt das A'-Segment verkürzt wieder (das ist schon die "erste Schwalbe" der intonationellen Entfaltung), wobei hier statt einer "weiblichen Endung" der ersten zwei Segmente eine "männliche Endung" vorliegt, die Melodie bleibt auf einem "g" stehen. Dies ist eine überraschende Abbremsung - der erste Vorbote des künftigen Melodierückganges, ein "Abbruch ihrer Aufwärtsbewegung"³). Die Melodie geht in der Tat nicht höher als zum erreichten "g". Weiter, es entsteht die folgende Sequenz von F-Motiven: ein Ton, zwei Töne, drei Töne. Als ob die Melodie nach der Stillsetzung allmählich zu Atem käme - ein besonders wichtiges Moment, das eine wesentliche Facette des lyrischen Charakters des Themas

enthüllt. Schließlich, das letzte dreitönige F-Motiv bildet eine "weibliche Endung". Es erweist sich also, daß die Melodie eine Reprisenstruktur auf der Mikroebene hat. Es ist vollkommen klar, daß alle diese Erscheinungen einfach unsichtbar sind, wenn man die Melodie mit Hilfe größerer Einheiten als F-Motiven analysiert. Hier kann man noch hinzufügen, daß das F-Motiv die kleinste bis heute bekannte Einheit mit variabler Länge ist und daß die F-Motiv-Segmentierung für unsere Wahrnehmung deswegen so ungewohnt ist, weil sie auf Gesetzmäßigkeiten der musikalischen Mikrogliederung beruht, die uns meistens nicht bewußt werden (es ist charakteristisch, daß die Erfüllung des Zipf-Mandelbrotschen Gesetzes im großen Maße mit der unterbewußten Kontrolle der Elementenwiederholung auf der Seite des Hörers oder des Autors zusammenhängt, vgl. Boroda, Orlov 1978).

Zu J. Wildgruber

Ich stimme mit Wildgruber voll überein, daß Wiederholungen die wichtigste Regularität der Konstruktion musikalischer Zeichen bilden. Eben von dieser Tatsache bin ich ausgegangen, als ich die Organisation der Wiederholbarkeit der Mikroelemente im musikalischen Text zu analysieren anfang. Daraus folgt jedoch keineswegs, daß man die Rolle der nur einmal im musikalischen Text vorkommenden Elemente - hier einmaliger F-Motive - unterschätzen soll: ihre Anzahl beträgt ungefähr die Hälfte des gesamten F-Motiv Inventars (ebenso im literarischen Text: die Zahl der einmaligen Wörter umfaßt etwa die Hälfte des Textwortschatzes). Weiter, die erhaltenen Resultate beweisen, daß einmalige F-Motive keineswegs einen "musikalischen Leerlauf" bilden: wie ein Blick in die musikalischen F-motivischen Häufigkeitswörterbücher, die ich für jeden untersuchten musikalischen Text erhielt (in meinem Aufsatz konnte ich ein solches Wörterbuch leider nicht einmal teilweise präsentieren)⁴), zeigt, sind gerade die einmaligen F-Motive vom intonations-melodischen Standpunkt die interessantesten F-Motive, da sie die ursprünglichen Tonfolgen enthalten. Die häufigen F-Motive sind im Intonationsplan viel "ärmer", sie bilden die sogenannten "allgemeinen Formen der musikalischen Bewegung", den "Hintergrund". Außerdem zeigt eben die Analyse des Zipf-Mandel-

brotschen Gesetzes, daß die Bedeutung einer solchen Charakteristik wie Textlänge für die Organisation der F-Motive oder Wortwiederholungen gerade damit zusammenhängt, daß in jedem Text eine große Anzahl nicht wiederholter F-Motive (Wörter) vorhanden ist. Es ist noch zu bemerken, daß vom informationstheoretischen Gesichtspunkt eben die seltensten Elemente die größte Menge der Information tragen. Schließlich, für ein normales Kunstwerk - speziell ein musikalisches - ist die ständige Zunahme neuer Elemente unvermeidlich, und diese Zunahme wird, wie die Rechnungen zeigen, vor allem durch Einfügung seltener Wörter (größtenteils einmaliger Wörter) gewährleistet.

Sehr interessant ist Wildgrubers Bemerkung, daß in der modernen Musik (speziell der mit Reihentechnik) die Identifizierung der F-Motive problematisch wird. Dies ist eine Anregung zur Untersuchung der Organisation der F-Motiv-Wiederholungen bei unterschiedlichen Identifikationskriterien: beispielsweise bei Kriterien, die nicht nur die exakte (strenge), sondern auch die modifizierte (freie) Sequenztransposition berücksichtigen usw. Eine solche Untersuchung wäre auch deswegen sinnvoll, weil man dadurch das für jedes musikalische Werk optimale Wiederholungskriterium auf der Mikroebene entdecken könnte (hier wird man die Organisationsform allerdings postulieren müssen, z.B. als die Erfüllung des Zipf-Mandelbrotschen Gesetzes). Und mit dem Problem des Wiederholungskriteriums sind wichtige Bedingungen der Rezeption des musikalischen Werkes verbunden. Vielleicht liegt hier einer der Schlüssel zur Problematik der Komplexität der Rezeption "moderner" Musik.

Sehr wichtig ist auch die andere Seite der Bemerkung von Wildgruber: das Problem an sich, F-Motive in Werken zu segmentieren, die mit einer Kompositionstechnik geschrieben wurden, die sich von der Technik des XVIII-XIX. Jhdt. sehr unterscheidet (z.B. die Reihentechnik). Dies ist ein interessantes Problem, und ich habe es bei der Beurteilung der Perspektiven der Untersuchung von Mikrowiederholungen in musikalischen Texten in meiner Dissertation berührt (vgl. Boroda 1979). Auf der einen Seite spielen in vielen Werken, z.B. bei Webern, ganz kleine Gebilde, "Mikromotive", denen das F-Motiv einigermaßen nahe steht, eine sehr wichtige Rolle. Auf der anderen Seite ist die Lösung des Problems einer

"effektiven" Elementareinheit in der Musik nur experimentell möglich: man muß zählen, zählen und zählen. In diesem Sinne hat die Feststellung irgendwelcher invarianten Gesetzmäßigkeiten auf der Ebene der gegebenen Einheit noch einen anderen Wert: Man kann diese Gesetzmäßigkeiten als einen Ausgangspunkt betrachten, von dem aus man die Einheiten und die Kriterien ihrer Unterscheidung (Wiederholungskriterien) variiert, und hinter den Veränderungen die Formen der Gesetzmäßigkeit beobachtet. Wegen der wichtigen Rolle der Wiederholung in der Musik bietet die Analyse der Wiederholungsstrukturen in diesem Sinne besonders reichhaltige Möglichkeiten.

Bezüglich der von Wildgruber erwähnten Divergenz zwischen dem F-Motiv und der "motivischen Gestalt" kann man schließlich nur wiederholen: das F-Motiv wurde nicht als ein volles oder partielles "formales Äquivalent" des Motivs definiert (in diesem Sinne ist die Bezeichnung meiner Einheit als "formales Motiv" nicht sehr gut gewählt worden); außerdem ist das F-Motiv eine betont kleine Einheit.

Zu W.-D. Schäfer

(a) Zu der Bemerkung von Schäfer, daß bei den Abweichungen der empirischen Werte von n , n_1 , n_2 , n_3 von den theoretischen Prognosen sich die historische Entwicklung der Musik vom Barock bis zum 20. Jhdt. keineswegs spiegelt: Dies ist in der Tat so. Aber es ist auch natürlich. Denn meine Aufgabe bestand in der Suche nach metastilistischen Gesetzmäßigkeiten der Wiederholungsorganisation, in der Suche nach Charakteristika, die trotz der starken Entwicklung der europäischen Musik im Bereich der Harmonie, Melodie, Form usw. invariant sind. Um solche Gesetzmäßigkeiten zu finden, war es nötig, eine solche musikalische Einheit zu finden, die sich an sehr stabile Gesetzmäßigkeiten der Melodieorganisation anlehnt, an Gesetzmäßigkeiten, die sich sozusagen im Fundament der taktorganisierten Musik, die ich untersuchte, befinden. Gerade das F-Motiv lehnt sich an solche Gesetzmäßigkeiten an, und von diesem Standpunkt aus hat es schon an sich metastilistische Allgemeinheit. Das F-Motiv als Einheit spielt bei den

erzielten allgemeinen Resultaten eine sehr bedeutende Rolle, auch wenn sie vorläufig noch nicht exakt einzuschätzen ist.

Bezüglich der Bemerkung von Schäfer, daß das F-Motiv eine Art "Fiktion" zu sein scheint, muß ich sagen, daß das F-Motiv, wie die Untersuchung zeigt, als eine aktuelle elementare Einheit in unterschiedlichen musikalischen Stilrichtungen erscheint (bezüglich des Zusammenhanges des F-Motivs mit den feinen Besonderheiten der melodischen Phrasierung, der Mikrostruktur der Melodie, vgl. meine Antworten an Fischer). Man kann auch nicht sagen, daß das F-Motiv beispielsweise bei Prokofiev weniger fruchtbar wäre als bei Bach oder Tartini. Es wäre, letzten Endes, ratsam, ersthafte Untersuchungen durchzuführen, aber dies ist keineswegs so einfach: schon aus dem Grunde nicht, weil ein "Maßstab" für eine semantische Einheit fehlt (das Motiv und das Teilmotiv erfüllen diese Aufgabe leider nicht - vgl. meine Antworten zu Fischer).

(b) Die von Schäfer angeführten Resultate sind äußerst interessant (leider gibt Schäfer keine exakte Definition seiner Zähl-einheit, was mein Verständnis seiner Daten etwas erschwert). Jedoch ist die Tatsache, daß diese Resultate irgendwie historisch interpretierbar sind, überhaupt kein Argument für die Einwände gegen "ahistorische" Resultate, die mit Hilfe des F-Motivs erhalten wurden. Die Einheit, mit der Schäfer die Zählungen durchgeführt hat, muß nämlich historisch abhängig sein (Schäfer selbst schreibt darüber), die Beziehung zum timbre, d.h. auch die Behandlung des "Klangrepertoires" durch den Komponisten änderte sich mit der Zeit. In diesem Sinne ist "Klangrepertoire" keineswegs dem F-Motiv äquivalent.

Schäfers Bemerkung berührt jedoch das äußerst wichtige Problem der "natürlichen Einheiten" und der Durchzählung des musikalischen Werkes aufgrund unterschiedlicher Einheiten - ein Problem, das einen "Engpass" der quantitativen Analyse der Musik überhaupt darstellt. Zweifellos erlaubt eine solche mit Hilfe unterschiedlicher Einheiten durchgeführte vielseitige Analyse, eine Reihe wichtiger Zusammenhänge im musikalischen Werk zu entdecken, diverse in ihm ruhende Gesetzmäßigkeiten zu einem allgemeinen Bild zu verknüpfen und auf dieser Basis ein Modell des musikalischen Werkes zu konstruieren. Urteilt man jedoch nach dem heu-

tigen Stand der quantitativen Musikanalyse, so liegt dies in einer fernen Zukunft.

(c) Bezüglich Schäfers Bemerkung, daß die Schätzungen von n_m einer erheblichen Verbesserung bedürfen: Der von mir entdeckte Effekt der "zweifachen Aufführung des F-Motivs" zeigt, daß im Bereich niedriger Häufigkeiten (d.h. im Bereich der selten vorkommenden F-Motive des musikalischen Textes) gegenseitige Kompensationen, Umgruppierungen usw. durchgeführt werden. Unter diesen Bedingungen Korrekturkoeffizienten der Schätzungen von n_m einzuführen, würde bedeuten, daß man "diesen Effekt durch Formeln zum Gesetz macht", anstatt daß man ihn als eine bestimmte Abweichung der Wiederholungsstruktur im musikalischen Text vom Zipf-Mandelbrotschen Gesetz betrachtet, d.h. als eine (eventuelle) Eigenart der Organisation des musikalischen Textes. Es ist zweifelhaft, ob es sinnvoll ist, mit einem Effekt, dessen Natur und Erzeugungsmechanismen noch nicht geklärt sind, so schnell "fertig zu werden" (die in meiner Arbeit angeführten Daten reichen in diesem Sinne nicht für endgültige Folgerungen aus). Auf der anderen Seite, wenn man das Problem ernsthaft betrachtet, muß die Korrektur der Schätzungen von n_m mit der Korrektur des ganzen Modells der Wiederholungsstruktur und der Form des Zipf-Mandelbrotschen Gesetzes, zu der dieses Modell führt, zusammenhängen. Eine solche Arbeit kann jedoch nur unter der Bedingung angefangen werden, daß man ernsthafte, sehr gewichtige Beweise für ihre Unerläßlichkeit vorbringt. Auf jeden Fall, wie ich es sehe, braucht man vorläufig (auch wenn nur vorläufig) nicht nach einer exakten Beschreibung beliebiger Abweichungen von festgelegten Strukturen mit Hilfe von Formeln zu streben: Sie verhindern die aufmerksame Ausschau nach spezifischen Texteffekten (die vielleicht die Musik von der Literatur unterscheiden usw.). Man soll lieber vorsichtig nach der Ablösung des Modells streben.

Ich bedanke mich bei allen Diskussionsteilnehmern für die aufmerksame Lektüre meines Aufsatzes und ihre Bemerkungen, die mir erlaubt haben, meine Ansichten ausführlicher darzustellen, als es im Aufsatz möglich war.

Anmerkungen

1. Unter dem Titel Orlov & Boroda & Nadarejšvili, Sprache, Text, Kunst. Quantitative Analyse (1980).
2. Um Mißverständnisse zu vermeiden, möchte ich folgendes bemerken: wenn ich dieses Beispiel bringe und Fischer widerspreche, so will ich damit keineswegs sagen, daß keine Melodie in Segmente fixierter Länge zerlegt werden kann (z.B. in Segmente von n Tönen): eine Reihe von Melodien, in denen eine monotone rhythmische Bewegung herrscht, Töne gleicher Länge - z.B. viele Etüden, Stücke des Typs von moto perpetuo - erlauben ohne weiteres eine solche Zerlegung. Aber auch in solchen Fällen liegt das Problem darin, daß die Länge des für die Melodie natürlichen Segments mit fixierter Länge vor der Segmentierung, vor der Analyse dieser Melodie nicht festgelegt werden kann. Ein Segment, das in einer Melodie mit gleichlangen Tönen natürlich ist, kann in einer anderen Melodie ganz unnatürlich sein. Ganz große Schwierigkeiten ruft die Zerlegung der Melodie in Elemente mit fixierter Länge in dem Fall hervor, wenn in der Melodie Töne unterschiedlicher Längen benutzt werden: in dem Fall entsteht ein natürlich klingendes Segment wie durch Zufall, ganz unerwartet inmitten "gehaltsamer" Zerlegungen.
3. Dieser Begriff wurde aus Orlov (1980) übernommen, wo man eine interessante Analyse dieses Themas von Čajkovskij findet.
4. Vollständige F-Motiv-Häufigkeitswörterbücher einer Reihe musikalischer Werke, die in Tab. 1 aufgeführt sind, kann man in meiner Dissertation (Boroda 1979) finden. Ebenda findet man auch ausführliche Angaben über die Häufigkeitsstruktur musikalischer Werke auf der Ebene der F-Motive für alle Werke in der Tab. 1. Einige Angaben über solche F-Motiv-Häufigkeitswörterbücher und ihre Analyse findet man in Orlov (1980).

LITERATUR

- ARAPOV, M.V., EFIMOVA, E.N., ŠREJDER, JU.A., O smysle rangovyh raspredelenij. Naučno-Tekničeskaja informacija, 2, 1975/1, 9-20.
- ARAPOV, M.V., ŠREJDER, JU.A., Klassifikacija i rangovye raspredelenija. In: Naučno-Tekničeskaja informacija 2, 1977/ 11-12, 15-21.
- BORODA, M.G., O melodičeskoj elementarnoj edinice. In: Gosovskij, V.L. (Hrsg.), Materialy Pervogo Vsesojuznogo seminara po mašinnyh aspektam algoritmičeskogo formalizovannogo analiza muzykal'nych tekstov. Erevan, Izdatel'stvo AN Armjanskij SSR 1977, 112-120.
- BORODA, M.G., Principy organizacii povtorov na mikrourovne muzykal'nogo teksta. Diss. Tbilisi 1979.
- BORODA, M.G., ORLOV, JU.K., O nekotorych psihologičeskich aspektach količestvennoj organizacii chudožestvennyh tekstov. In: Bassin, F., Prangišvili, V., Serozija, A. (Hrsg.), Bessoznatelnoe. Priroda, funkcii, metody issledovanija. Tbilisi, Mecniereba 1978, Bd. III, 302-309.
- GROVE's dictionary of music and musicians. London, Macmillian 1972, Vol. 3.
- LOBE, K., Musikalnij katehisis. Moskva, Yurgenson 1887.
- MANDELBROT, B., O rekurrentnom kodirovanii, ograničivajuščem vlijanie pomech. In: Sifozov, W.J. (Hrsg.) Teorija peredači Soobščenij. Perevod s anglijskogo, Moskva, Izdatel'stvo inostrannoju literatury 1957.
- MAZEL', L.A., Stroenie muzykal'nych proizvedenij. Moskva, Muzyka 1978.
- MAZEL', L.A., CUKERMAN, V.A., Analiz muzykal'nych proizvedenij. Moskva, Muzyka 1967.
- MUZYKAL'NAJA ENCIKLOPEDIJA. Moskva, Sovetskaja enciklopedija 1976.
- ORLOV, JU.K., Obobščennyj zakon Cipfa-Mandel'brot'a i častotnaja struktura informacionnyh edinic različnyh urovnej. In: Guseva, E.K. (Hrsg.) Vyčislitel'naja lingvistika. Moskva, Nauka 1976, 180-202.

- ORLOV, JU.K., Model' častotnoj struktury leksiki. In: Andruš-čenko, W.M. (Hrsg.) Issledovanija v oblasti vyčislitel'noj lingvistiki i lingvostatistiki. Moskva, Izdatel'stvo Moskovskogo universiteta 1978, 59-118.
- ORLOV, JU.K., Nevidimaja garmonija. In: Čislo i mysl'. Moskva, Znanie 1980 (im Druck).
- PRICE, D. de S., A general theory of bibliometric and other cumulative advantage processes. Journal of the American Society for Information Science 27 (5), 1976, 292-306.
- RIEMANN, H., Musik Lexikon. Berlin, Hesse 1929.
- RIEMANN, H., Musik Lexikon. Mainz, Schott's Söhne 1967.
- SEEGER, H., Musiklexikon in zwei Bänden. Leipzig, VEB, Deutscher Verlag für Musik 1966, Bd. II.
- TJULIN, JU.N. (Hrsg.) Muzykal'naja forma. Moskva, Muzyka 1974.

THE BUTTERFLY REVISITED

A RE-ANALYSIS OF DEESE'S STUDY 'ON THE STRUCTURE OF ASSOCIATIVE MEANING'

W. Marx, G. Strube, München

Abstract

Deese's 1962 study of butterfly associates has been criticised both for its interpretation of results and its method, factor analysis being judged inadequate for use with overlap coefficients. This led to re-analysis of Deese's data by means of nonmetric multidimensional scaling. The resulting two-dimensional structure showed clusters identifiable with four of Deese's six factors.

In his 1962 paper, "On the structure of associative meaning", Deese performed the first experimental analysis of a linguistic field. This study has been recognized as a classic because, as Hörmann (1967) puts it in his extensive review (now deleted in the second edition), it challenges subjective, 'armchair' taxonomies of vocabulary and concentrates on language as spoken by a group of people.

Criticism, then, is not intended to question the importance of this pioneering work, but is concerned with problems of method. Deese's study, it has been said (Strube, 1979:167) "is handicapped by the use of statistical methods hardly suited to the task, more appropriate methods not being available at that time".

Deese proceeded as follows: (1) A word, viz. butterfly, and 18 of its most common associates were taken as stimuli for a word-association task. (2) The associative "overlap" between each two of the response distributions was computed, resulting in a matrix of overlap coefficients (see table 1). (3) Factor analysis of the overlap matrix, resulting in a six-factor solution "of convenient simple structure" (Deese 1962, 168). The rotated factor loadings are shown in table 2. (4) The first four factors were interpreted in terms of a hierarchical classification scheme.

Factor analysis is usually based on correlation matrices which, being symmetrical, resemble a matrix of overlap coefficients (OC). Overlap coefficients, however, make use of the frequency of common responses to each of two stimulus words, OC being the ratio of common responses to total number of responses given. Thus overlap coefficients differ fundamentally from correlation-like similarity coefficients, the latter being scalar products. But factor analysis, at least in its usual form as employed by Deese, needs similarity measurement in terms of scalar products, therefore overlap coefficients are not suited for factor analysis.

Fortunately, techniques of nonmetric multidimensional scaling (NMDS) have been developed in the meantime, thus resolving the dilemma (NMDS is not restricted to similarity coefficients of the scalar product type). This stipulated the question of whether Deese's factorial structure would conform to the structure obtained in a re-analysis of his data by means of a NMDS technique.*

Method and results

Deese's matrix of overlap coefficients was used as input to the MINISSA program by Roskam and Lingoes (1975 version). Solutions in 2 and 3 dimensions, each using Minkowski coefficients of $r = 1$ (city-block metric), 2 (Euclidean), and 8 (for approximation of supremum metric), were computed. Minimum stress was obtained for Euclidean metric with a 3-dimensional solution (stress $\hat{d} = .06$, stress $d^* = .07$). Stress was but slightly greater with two dimensions (.09 and .12, respectively), so the two-dimensional solution was chosen. Table 3 shows the coordinates, while figure 1 gives a graphical representation of the structure.

Discussion

Comparing the NMDS structure of fig. 1 with Deese's factors (table 2), the first four of these can be recognized easily in the cluster structure produced by NMDS. There are, of course, some

* Readers not familiar with the method will find a short introduction to NMDS in Strube (in press).

minor discrepancies, viz. items 13 (summer) and 17 (nature) are isolated, and item 19 (butterfly) goes with the items that have positive loadings on Deese's third factor (III+), which indeed makes better company than its original III-grouping (see below). These details excepted, factor analysis and NMDS yield practically equal solutions, though Deese's factors V and VI fail to get represented in the NMDS structure.

It is the same factors (V and VI) that give rise to problems in the original interpretation by Deese, the hierarchical structure of which is maimed by the fact that both these factors run wildly across the grouping as established by means of factors I through IV. Interestingly enough, this issue is not dealt with in Deese's discussion (1962, 169), not even the factors V and VI as such.

Deese does not give information on his criterion for determining the number of factors, but we may suspect him of having extracted two superfluous factors, perhaps in order to get some 43 instead of 32 percent of variance explained (data computed from the matrix of factor loadings given in table 2). "More factors than needed, account for a less than satisfying portion of total variance, and the strategy proposed for interpretation has failed." (Strube 1979.)

With the exception of amount of variance explained, all these difficulties can be overcome by dispensing with factors V and VI. We hold that a factor solution which lends itself readily to interpretation (if possible, in terms of an already established theory) is to be preferred to one obtained 'blind', i.e. according to statistical criteria only. Purely statistical criteria (e.g., for number of factors (dimensions) should not be preferred to criteria of interpretability (cf. Gigerenzer & Strube, 1978). If this argument is accepted, Deese's factorial and our NMDS structure may be taken as equivalent. This result suggests that factor analysis is not too much affected by improper similarity measures.

There remain some aspects of content to be discussed. Deese's factors I through IV define the following two groups, divided into two subgroups each (see Deese, 1962, 169; factors are indicated by Roman numbers, positive and negative loadings by + and - signs,

respectively):

I+: ANIMATE III+: wing, bird, fly, bees, (nature)
 III-: butterfly, cocoon, insect, moth, bug
 II+: NONANIMATE IV+: yellow, blue, sky, color
 IV-: garden, flower, spring, summer, sunshine

The classification resulting from NMDS, while replicating the above grouping in general, places nature and summer apart from the clusters and butterfly in the (I+, III+) group. Nature, in fact, is isolated in the factorial grouping, too, its loading on factor I being nearly zero, so its NMDS position is rather clarifying on the point that nature as such is not an animal.

Factor analysis places summer firmly in the (II+, IV-) group, with loadings of .31 and -.34, resp. By NMDS, however, it is placed isolated, the nearest word being spring, i.e. the only word sharing a substantive overlap with summer.

Note that in the "nonanimate" group, the subdivision by factor IV loadings is questioned by NMDS: it seems possible but not essential to subdivide this rather heterogeneous group, the organization of which is obviously not determined by logic (nor is flower inanimate), but by contiguity, i.e. through the experience of man in everyday life.

In the "animate" group, butterfly has broken its ties with cocoon (and all the noxious and ugly bugs, the moth, and insects), and has taken sides with the birds and bees that have wings and fly. This is due to the fact that NMDS, as well as factor analysis, and contrary to clustering techniques proper, recognizes distances between items on the whole, thus sometimes slightly neglecting the immediate surroundings of an item (as here does NMDS with the strong butterfly-cocoon proximity, or Deese's factor analysis has done in the case of summer). It is a minor point, however, since a subdivision of the NMDS cluster seems not essential.

Deese interpreted the results of his factor analysis as indicating an "associative thesaurus" (1962, 169), thus suggesting a hierarchical classification by semantic features in spite of his cautious words "that associative meaning is not logical" (ibid.). NMDS, however, produces what might be called a non-hierarchical cluster structure (i.e., localization in a multidimensional space),

organized by the principles of similarity and of contiguity in everyday experience, principles that since Locke and Hume are recognized as 'laws' governing association.

References

- DEESE, J., On the structure of associative meaning. Psychological Review, 1962, 69, 161-175.
- GINGERENZER, G., STRUBE, G., Zur Revision der üblichen Anwendung dimensionsanalytischer Verfahren. Zeitschrift für Entwicklungspsychologie und Pädagogische Psychologie, 1978, 10, 75-86.
- HÖRMANN, H., Psychologie und Sprache. Berlin, Heidelberg, New York: Springer, 1967.
- STRUBE, G., Associative meaning and its analysis. In: Grotjahn, R. (Ed.), Glottometrika 2, Bochum, Brockmeyer 1979, 158-197.

Table 1. Overlap Coefficients for Common Associates, from Deese (1962, 167) (Decimals omitted)

Stimulus words	Stimulus words																		
	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19
1 MOTH	100	12	12	12	11	02	00	05	11	00	00	02	02	05	01	01	01	01	15
2 INSECT	100	09	09	17	01	01	01	33	10	01	01	03	00	00	00	00	01	00	12
3 WIND	100	44	19	00	00	00	00	03	02	00	00	10	00	00	00	00	03	00	13
4 BIRD	100	21	01	00	03	02	01	01	00	01	01	10	00	01	00	01	05	00	12
5 FLY	100	01	01	08	06	01	01	08	06	01	02	06	00	03	00	02	04	00	11
6 YELLOW	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
7 FLOWER	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
8 BUG	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
9 COCOON	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
10 COLOR	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
11 BLUE	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
12 BEES	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
13 SUMMER	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
14 SUNSHINE	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
15 GARDEN	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
16 SKY	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
17 NATURE	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
18 SPRING	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06
19 BUTTERFLY	100	07	00	00	17	02	02	02	00	03	07	02	01	06	18	02	06	02	06

Table 2. Rotated Centroid Factor Loadings from Deese (1962, 169) (Decimals omitted)

Words	Factors					
	I	II	III	IV	V	VI
Moth	44	03	-27	-01	-03	-32
Insect	50	01	-33	01	-34	11
Wing	52	01	45	01	29	-07
Bird	52	02	46	01	29	-07
Fly	48	03	32	01	-28	-03
Yellow	01	44	-03	34	-32	-02
Flower	01	39	-03	-32	03	44
Bug	41	01	-34	00	-14	37
Cocoon	40	01	-35	00	25	02
Color	-02	42	-04	44	04	-04
Blue	-02	57	-04	52	23	-04
Bees	36	04	34	-02	-30	00
Summer	-01	31	-03	-34	-02	-34
Sunshine	02	37	-04	-33	-03	-35
Garden	00	35	-02	-34	-03	44
Sky	-01	41	-03	43	38	-07
Nature	04	31	29	-02	01	34
Spring	-01	35	-03	-37	-02	-36
Butterfly	48	06	-29	-01	26	01

Table 3. Coordinates of the Two-Dimensional NMDS Structure

Stimulus words	x ₁	x ₂
Moth	-.8094	-.0402
Insect	-.9471	-.3742
Wing	-.5873	-.6450
Bird	-.3805	-.6023
Fly	-.4410	-.4312
Yellow	.8784	.1100
Flower	.5356	.3177
Bug	-1.1096	-.3895
Cocoon	-.9631	.0954
Color	1.4350	-.3035
Blue	1.1818	-.1613
Bees	-.1599	-.1304
Summer	-.7070	1.5254
Sunshine	.5994	.6987
Garden	.2588	.7177
Sky	1.3114	-.5946
Nature	.0867	-.9161
Spring	.1311	1.1216
Butterfly	-.3133	.0017

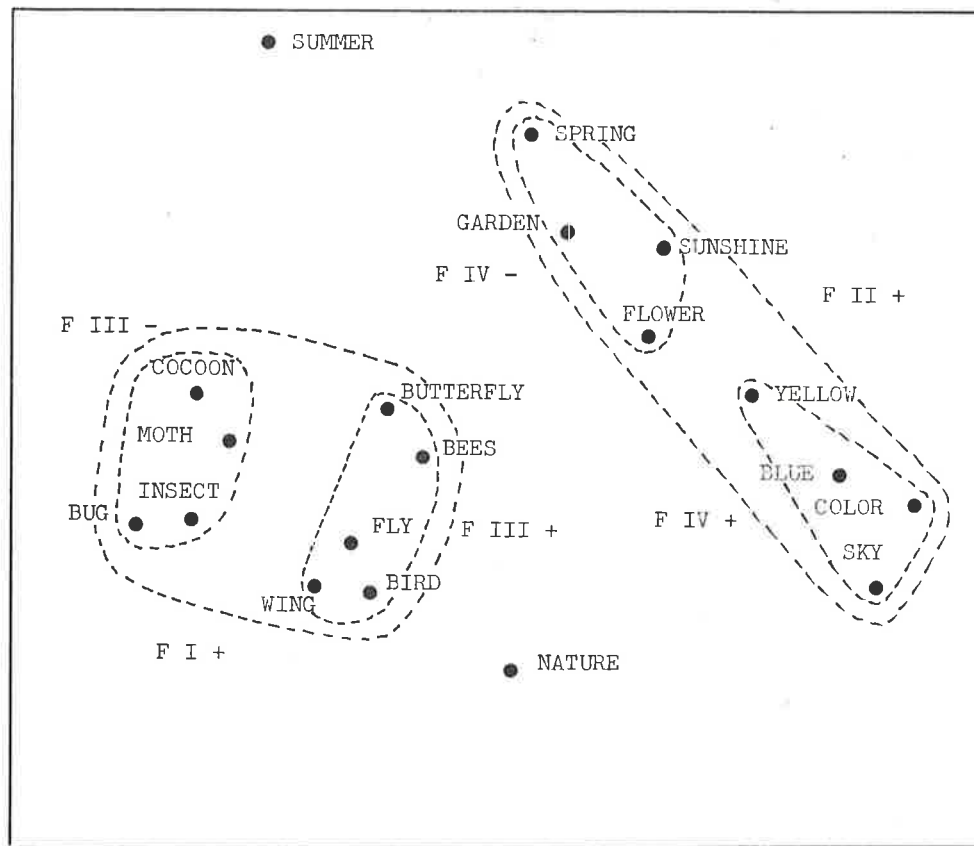


Figure 1.
Two-Dimensional NMDS Structure of Butterfly and its Associates

GEDÄCHTNIS UND ERKENNEN

EIN BEITRAG ZUR SEMANTISCHEN THEORIE DER ERKENNUNGSPROZESSE

L. Kolman, Prag

1. Einführung
2. Gedächtnismodelle
 - 2.1 Das Erkennen: Die aktive und die passive Theorie
 - 2.2 Serielle und parallele Prozesse
 - 2.3 Die Gedächtnisstruktur
 - 2.4 Das Gedächtnis als Katalog
3. Semantische Strukturen
 - 3.1 Die Semiotik des Erkennens
 - 3.2 Semantische Strukturen
 - 3.3 Einige Anmerkungen
4. Messungen und Experimente
 - 4.1 Messung der semantischen Distanz
 - 4.2 Die Experimente
 - 4.2.1 Durchführung der semantischen Distanzmessung
 - 4.2.2 Freie Reproduktion und Wiedererkennen
 - 4.2.3 Paarassoziations-Lernen
 - 4.3 Auswertung der semantischen Distanzmessung
 - 4.4 Semantische Distanz und Gedächtnisleistung
 - 4.4.1 Freie Reproduktion und Wiedererkennen
 - 4.4.2 Paarassoziations-Lernen
- 4.5 Diskussion
5. Der Zusammenhang einer hierarchischen Gedächtnisstruktur mit einem induktiv arbeitenden Erkennungsmodell
 - 5.1 Diskriminationsstärke von Erkennungsmodellen
 - 5.2 Induktives Erkennen
 - 5.3 Semantische Merkmale
 - 5.4 Möglichkeiten und Grenzen der hierarchischen Strukturen
6. Eine formale Theorie der semantischen Strukturen
 - 6.1 Ein Differenzmaß
 - 6.2 Semantische Strukturen
 - 6.3 Der Graph einer semantischen Struktur

Literatur

1. EINFÜHRUNG

Erkennen ist die Zuordnung dessen, das wir rezipieren, zu etwas, das wir bereits kennen, also "im Gedächtnis" haben. Damit sich eine solche Zuordnung vollziehen kann, muß das rezipierte Material irgendwie mit gespeichertem Material verglichen werden. Wäre das Gedächtnis ein ungeordneter Speicher, angefüllt mit den Repräsentanzen der früher wahrgenommenen Objekte, gleichsam Bildern, mit denen die aktuell rezipierten Reizmuster zu vergleichen sind, so wäre das Mustererkennen eine sehr zeit- aufwendige Aufgabe; denn die Zahl solcher Gedächtnisbilder muß sehr groß sein. Deswegen müssen wir annehmen, daß der Gedächtnisspeicher eine gewisse innere Struktur hat, die das Auffinden erleichtert.

Dieser Gedanke führt zur Konzeption eines katalogartig geordneten Gedächtnisses, vergleichbar dem Katalog einer Bibliothek. Einen solchen "Katalog" nennen wir im nachfolgenden Text eine Gedächtnisstruktur. Daß die Gedächtnisstruktur in den Erkennungsprozessen eine bedeutsame Rolle spielt, läßt sich bei einem Exkurs in die Theorie des Musteridentifizierens zeigen. Die gegenwärtigen Ansichten über Erkennungsprozesse sind übrigens mit mehreren denkbaren Strukturmodellen des Gedächtnisses vereinbar. Insofern man die Struktur des Gedächtnisses als eine Struktur der Charakteristika des gespeicherten Stoffs, der identifizierten Inhalte, ansehen kann, stellt sich ihre kognitive Funktion als ein Problem im Rahmen der Semiotik. Die semiotische Analyse führte zum Vorschlag von zwei alternativen semantischen Strukturmodellen des Gedächtnisses.

Um die beiden semantischen Strukturen genau zu definieren, wurden sie als mathematische Modelle dargestellt. Wie man zeigen kann, unterscheiden sich die beiden Strukturen hinsichtlich der in ihnen ablaufenden Aktualisierungsprozesse. Die semantischen Entfernungen der zu reproduzierenden Inhalte haben in

diesen Modellen unterschiedliche Auswirkungen auf die Reproduktion.

Da sich für die vorgeschlagenen Modellalternativen Unterschiede im Reproduktionsprozeß ableiten ließen, war es möglich, die Modelle in psycholinguistischen Experimenten zu testen, in denen zur Messung der semantischen Entfernung das von G. A. MILLER eingeführte Maß verwendet wurde. Die Resultate sprechen für eine hierarchische Gedächtnisordnung.

Schließlich werden die Ergebnisse aus der Sicht der Theorie der Mustererkennung diskutiert und zwar unter der Frage, welche Rolle die semantische Struktur des Gedächtnisses in diesen Prozessen spielt. Wir schlagen vor, das Erkennen als einen induktiven Prozeß zu sehen. Es werden einige Befunde angeführt, die über die Natur dieses Prozesses Aufschluß geben.

2. GEDÄCHTNISMODELLE

Vom Standpunkt des Assoziationismus aus ist der Begriff des Gedächtnisses in der Psychologie überflüssig. Man sieht hier das Behalten als die Ausbildung neuer Verbindungen zwischen Aktivitätszentren im Nervensystem. Die kognitive Lerntheorie dagegen, die sich mit der Erforschung der Identifikation und des Begriffserwerbs befaßt, wobei sie in den Begriffen die koordinierende und steuernde Komponente des menschlichen Verhaltens sieht, muß offensichtlich für das Behalten der Begriffe eine gewisse Gedächtnisstruktur annehmen. Die Auseinandersetzung zwischen dem assoziationistischen und dem kognitiven Ansatz in der Lerntheorie überträgt sich auf das Gebiet des Mustererkennens und nimmt dort die Form der Auseinandersetzung zwischen einer passiven und einer aktiven Theorie des Erkennens an.

Es wird angenommen, daß Muster an bestimmten Zeichen, Charakteristika, erkannt werden. Parallelmodelle des Erkennungsprozesses nehmen an, daß diese Charakteristika simultan bestimmt werden. Bei seriellen Modellen dagegen wird die Annahme einer sukzessiven Ermittlung der Charakteristika gemacht.

Im Zusammenhang mit der Besprechung der erwähnten Auseinandersetzungen werden im folgenden zwei Gedächtnismodelle formuliert, die zugleich als Komponenten von Modellen des Erkennungsprozesses gesehen werden. Das eine Erkennungsmodell ist extrem passiv. Es wird gezeigt werden, daß es unzulänglich ist. Das andere Modell entspricht sowohl einer aktiven als auch einer gemäßigt formulierten passiven Theorie. Auf der Basis dieses zweiten Erkennungsmodells sollen die eigentlichen Modelle des Wiedererkennungsgedächtnisses erörtert werden.

Zum Schluß dieses Kapitels diskutieren wir die Befunde aus Untersuchungen, welche zur experimentellen Entscheidung darüber beitragen sollten, ob das menschliche Erkennen wie ein paralleler oder wie ein serieller Prozeß beschaffen ist. Die Befunde aus diesen Untersuchungen widersprechen sich. Manchmal ist sogar ihre Interpretation strittig. Diese Widersprüche ergeben sich, so scheint es nach allem, aus einer unzureichenden Definition dessen, was unter Charakteristiken zu verstehen ist.

2.1 Das Erkennen: Die aktive und die passive Theorie

Die aktive und die passive Auffassung entsprechen zwei grundlegenden Erkennungsmodellen, die ROSENBLATT (1958) aufgestellt hat. Im aktiven Modell ist es notwendig, "... daß die sensorische Information in kodierter oder bildlicher Form gespeichert wird, derart daß es eine eindeutige Abbildung zwischen den Reizen und den gespeicherten Mustern gibt". Im passiven Modell werden nicht Bilder aufgezeichnet, sondern "... das Zentralnervensystem funktioniert wie ein komplizierter Schaltmechanismus, in dem das Behalten durch die Schaltung neuer Verbindungen oder Wege zwischen Aktivitätszentren zustande kommt". Das passive Modell entspricht der traditionellen assoziationalistischen Auffassung vom Prozeß des Erkennens, insbesondere der Mediationstheorie. Das aktive Modell dagegen zieht zur Reizidentifikation gespeicherte Muster heran. Zu diesen Mustern gehören offensichtlich sowohl BRUNER's Katego-

rien als auch HUNT's Konzepte (1962, 1966) u.a. Wie übrigens UHR (1973: 189) erwähnt, unterscheiden sich die Computerprogramme zur Simulation des Begriffserwerbs und der Begriffsentdeckung praktisch nicht von einer Reihe von Programmen zur Simulation des Mustererkennens.

Sowohl die aktive als auch die passive Theorie haben ihre spezifischen Probleme. Die aktive Theorie muß zur Erklärung des menschlichen Erkennens einen Prozeß annehmen, der imstande ist, mit hoher Effektivität und Geschwindigkeit zu jedem sensorischen Reizmuster im Gedächtnis eine passende Repräsentanz aufzufinden. Diese Notwendigkeit veranlaßte LIBERMAN et al. (1964, 1967) zu der extremen Behauptung, das was wir erkennen, wenn wir jemanden sprechen hören, sei unsere eigene verdeckte Artikulation. Die passive Theorie kann liberaler formuliert werden, als sie ursprünglich von ROSENBLATT formuliert worden ist. Nichtsdestoweniger ist sie mit ernststen Problemen konfrontiert, wenn sie sich mit der hohen Variabilität der zu erkennenden Muster, wie sie beispielsweise beim Identifizieren handgeschriebener Buchstaben üblich ist, auseinandersetzen muß. Dies hat UTTLEY's (1958, 1966) gemäß der passiven Auffassung konstruiertes Modell demonstriert. Wenigstens in seiner älteren Version müßte es offensichtlich praktisch unendlich viele Neuronen enthalten, um das menschliche Erkennen modellieren zu können (UHR 1963).

Wenn wir das Gesicht eines Bekannten unter lauter fremden Gesichtern erkennen, so deutet die aktive Theorie das folgendermaßen: Wir vergleichen die Charakteristika der wahrgenommenen Gesichter mit denen der bekannten, und falls ein Reizmuster mit einer Gedächtnisrepräsentanz übereinstimmt, erkennen wir den Bekannten wieder. Im Gegensatz dazu nimmt die passive Theorie an, daß sich das Wiedererkennen des bekannten Gesichts unmittelbar mit der Bestimmung seiner Charakteristika einstellt. So werden uns z.B. eine spitze Nase, schmale Augenbrauen usw. direkt zu der Feststellung "Das ist Hans" veranlassen, dagegen eine stumpfe Nase, breite Augenbrauen (und weitere Merkmale) zu der Erkenntnis "Das ist Josef". Beide Ansätze nehmen übereinstimmend an, daß vor dem

Identifizieren eines Gebildes seine Charakteristika bestimmt werden müssen. Sie stimmen ferner überein in der Annahme einer bestimmten Gedächtnisrepräsentation dessen, was erkannt werden kann. Aber nach den aktiven Auffassungen sind im Gedächtnis Repräsentanzen gespeichert, wohingegen es nach den passiven Auffassungen aus Verschaltungen von Aktivitätszentren besteht. Bei seiner Formulierung des passiven Modells geht ROSENBLATT offensichtlich von der bekannten assoziationalistischen Analogie mit einem Telefonnetz aus. In der Tat, nach der passiven Theorie wirkt ein Sinneseindruck ähnlich wie eine Telefonnummer, die einen Anrufenden (hier das Reizobjekt) über eine Reihe von Umschaltungen mit dem gewünschten Anschluß (hier der passenden Reaktion auf den Reiz) verbindet.

In dem nach der aktiven Theorie arbeitenden Modell ist die Situation nicht so grundsätzlich anders, wie es nach ROSENBLATT's Formulierung den Anschein hat. Die Analogie mit der Telefonnummer könnten wir auch hier ganz gut anwenden, wenn auch mit dem Unterschied, daß wir hier statt von der Reaktion auf einen Reiz von der Abbildung oder Repräsentation des Reizes zu sprechen hätten. Wenn die passive Theorie liberal formuliert wird, d.h. wenn sie die bedingten Verbindungen nicht für gegenseitig unabhängige, isolierte Komponenten des Erkennungsprozesses hält, sind die beiden Theorien lediglich verschiedene Formulierungen derselben Grundvorstellungen über das Erkennen, Formulierungen in zwei unterschiedlichen Sprachen. Eine solche Doppelfassung führt zwar gelegentlich zu größerer Klarheit, hier aber wurde das Problem dadurch eher verdunkelt. Es handelt sich um das für beide Theorien gleichermaßen kritische Problem, einen Prozeß zu finden, der es ermöglicht, den Übergang vom Reiz zur Reaktion bzw. zur Repräsentanz mit hinreichender Effektivität, d.h. schnell genug und ohne allzu viele Fehler, zu vollziehen.

2.2 Serielle und parallele Prozesse

In den Theorien des Erkennens wird allgemein angenommen, daß ein sensorisches Reizmuster ein Ensemble von Charakteri-

stika ist, aufgrund derer der Stimulus identifiziert wird. Beim Identifizieren eines Objekts werden an ihm offenbar einzelne Charakteristika ermittelt. Von diesen hängt es ab, als was das Objekt wahrgenommen wird.

Zum Aufsuchen der Charakteristika eines Wahrnehmungsobjekts kann das Subjekt - wie übrigens jede Identifizierungseinrichtung - unterschiedliche Strategien benutzen. Diese Strategien kann man in zwei Gruppen einteilen, nämlich in serielle und parallele Strategien. Bei einer seriellen Strategie ermittelt das Subjekt ein Merkmal nach dem anderen. Bei der parallelen Suche dagegen werden alle Merkmale gleichzeitig bestimmt. Die Mehrzahl der Computerprogramme benutzt das serielle Suchen (vgl. z.B. BLEDSOE & BROWNING 1959). Nach der Argumentation von UHR (1973: 151) ist dies der Grund dafür, daß die Identifizierungsprogramme trotz der astronomischen Geschwindigkeit der Rechneroperationen so langsam sind. Ein Beispiel für serielles Suchen ist das Anwählen eines Telefonanschlusses, das wir im vorigen Abschnitt als Analogie verwendet haben. Die Ziffernfolge der Telefonnummer bestimmt sukzessiv, bei welchem Anschluß der Ruf ankommt.

Das serielle Prinzip wurde auch in einigen Computermodellen eingesetzt (z.B. HUNT et al. 1966). Es ist schneller, hat aber einen entscheidenden Nachteil. Es arbeitet nämlich weniger flexibel, und die Qualität der Identifizierungsleistung ist begrenzt durch die Zuverlässigkeit des schwächsten Gliedes in der Kette.

Mit dem Vergleich von Serien- und Parallelstrategien ist eine Reihe begrifflicher Unklarheiten verbunden. Die Auffassung CORCORAN's (1971) unterscheidet sich wesentlich von der UHR's (1973); CORCORAN hält die parallele Strategie für die schnellere. Er interpretiert Befunde aus Experimenten, die er mit seinen Mitarbeitern durchgeführt hat, als Beweis dafür, daß untrainierte Versuchspersonen nach der seriellen Strategie vorgehen (1967, 1971: 88). Nach ausreichendem Training entsprechen jedoch ihre Leistung der von einem Parallelmodell mit unbegrenzter Informationskapazität vorhergesagten. Dieser Schluß ist jedoch mehr als strittig, und UHR's Programm zur flexiblen Identifikation mit hierarchisch geordneten Charakterisatoren

(1973: 157) würde wohl ähnliche Resultate geben, ohne daß es einen Rechner mit unbegrenzter Kapazität verlangt.

2.3 Die Gedächtnisstruktur

Die passive Theorie des Erkennens und die assoziationalistische Theorie des Lernens tendieren beide dazu, die Rolle des Gedächtnisses einzuschränken. Im Gegensatz dazu wird die Funktion des Gedächtnisses von der aktiven Theorie sowie von der kognitiven Theorie des Lernens gerade betont. Wie wir aber in der vorangegangenen Diskussion gezeigt haben, folgen aus den Theorien gar nicht so unterschiedliche Ansichten über die Rolle des Gedächtnisses, wie es nach der ursprünglichen Formulierung ROSENBLATT's scheinen möchte. Der Schlüssel zur Erklärung dieser Tatsache liegt im klassischen assoziationalistischen Modell der bedingten Verbindung, die auch als S-R-Verbindung bezeichnet wird.

In diesem Modell sind die Grundbausteine, aus denen sich komplexere psychische Prozesse und schließlich das menschliche Verhalten im Ganzen aufbauen, die bedingten Verbindungen. Daß wir eine bedingte Verbindung scheinbar in Isolation untersuchen können, führt leicht zu der Ansicht, das menschliche Verhalten werde auf die allereinfachste Weise aus bedingten Verbindungen erzeugt, nämlich durch das bloße Aneinanderlegen von vielen solchen Verbindungen. Diese, man könnte sagen, extrem assoziationalistische Auffassung entspricht dem passiven Modell von ROSENBLATT, das nach der parallelen Strategie konzipiert ist.

Das nach diesem Prinzip konstruierte Modell des menschlichen Verhaltens ist auf Abb. 1 dargestellt. In diesem Modell gibt es nichts außer einer bestimmten Anzahl bedingter Verbindungen, die - zumindest innerhalb des Systems - keinerlei Zusammenhang miteinander haben. Jede dieser Verbindungen hat ihren Eingang, d.h. eine Einrichtung zur Reizaufnahme, und ihren Ausgang, d.h. eine Einrichtung zur Erzeugung der passenden Antwort. In der Abbildung sind Eingang und Ausgang durch eine unterbrochene Linie verknüpft. Die Unterbrechung der Li-

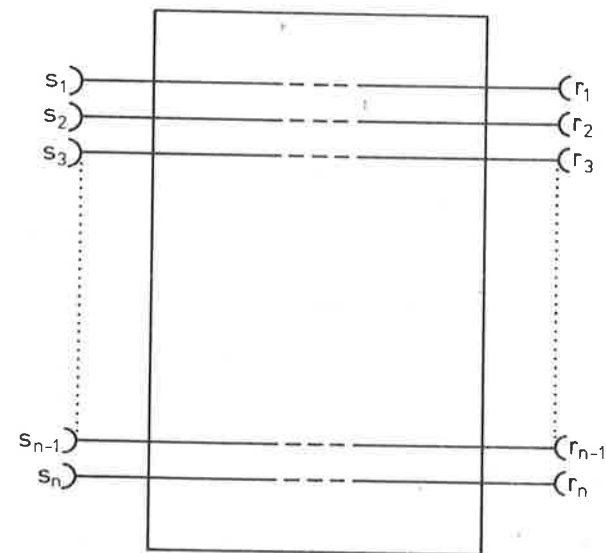


Abb. 1. Eine Kollektion isolierter Verbindungen.

$s_1, s_2, \dots, s_n$ sind Reize,

$r_1, r_2, \dots, r_n$ sind Reaktionen.

nie soll die Zeitweiligkeit der bedingten Verbindung darstellen. Die Paarungen von Eingängen mit Ausgängen sind Ergebnisse von Lernprozessen und können sich durch weiteres Lernen verändern. Wenn aber weiteres Lernen unterbleibt, so ändert sich das Verhalten des modellierten Systems nicht.

Sobald bei einem solchen System auf einen bestimmten Eingang ein passender Reiz einwirkt, wird sofort die Antwort ausgelöst. Das gilt auch dann, wenn mehrere Reize gleichzeitig auf das System einwirken. Dann werden alle zugehörigen Reaktionen gleichzeitig ausgelöst. Dank dieser Eigenschaft kann ein solches Modell auf jeden Reiz in der kürzest möglichen

Zeit reagieren. Wenn wir nun annehmen, daß das Modell Fehler machen kann - d.h. daß es mal in Abwesenheit eines adäquaten Reizes reagieren kann und manchmal auch auf einen adäquaten Reiz nicht reagiert - so wird die Fehlerhaftigkeit einer bedingten Verbindung die der übrigen keineswegs beeinflussen. Dieser Vorteil muß jedoch durch eine enorme Menge einzelner Verbindungen erkauft werden.

Wir haben versucht, das hier diskutierte Modell extrem passiv zu formulieren. Diese passive Auffassung ist - wie der Assoziationismus überhaupt - eine environmentalistische Auffassung. Sie sucht die Erklärung des menschlichen Verhaltens in Milieuwirkungen. Im Gegensatz dazu betont die aktive Theorie des Erkennens - ebenso wie die kognitive Lerntheorie - die Rolle der inneren Organisiertheit und Strukturiertheit des Zentralnervensystems, und zwar eine Strukturiertheit höherer Ordnung, als ein System isolierter bedingter Verbindungen sie impliziert. Die Modelle, mit denen wir uns jetzt beschäftigen werden, sind komplizierter strukturiert als eine Kollektion von isolierten bedingten Verbindungen. Ihre Struktur zeigt Abb. 2.

Das auf diesem Bild dargestellte Schema hat drei Blöcke: einen Eingang, eine zentrale Informationsverarbeitungsinstanz und einen Ausgang. Der mittlere Block besteht aus zwei Teilen, einem Gedächtnis und einer Entscheidungsinstanz. Der Gedächtnisblock gliedert sich in Charakterisatoren. Mit dem von UHR (1973) geprägten Ausdruck "Charakterisator" wird etwas bezeichnet, das wir auch Abbildung oder Repräsentanz eines Merkmals nennen könnten. Der Eingangsblock in Schema 2 entspricht den perzeptiven Prozessen. Am Eingang werden die Informationen aus der Umgebung, die Reize, in sensorische Reizmuster übersetzt. Diese werden dann mit den im Gedächtnis gespeicherten Charakterisatoren verglichen, wodurch ihre Merkmale bestimmt werden. Mit der Identifizierung seiner Merkmale ist das Reizmuster identifiziert, es muß nur noch über den passenden Output entschieden werden. Wenn wir noch annehmen, daß das Modell lernen kann - sei es, daß es seine Charakterisatoren vervollkommet und vermehrt, sei es, daß es sein Entscheidungsverhalten verbessert - so können wir es nicht nur als ein Modell

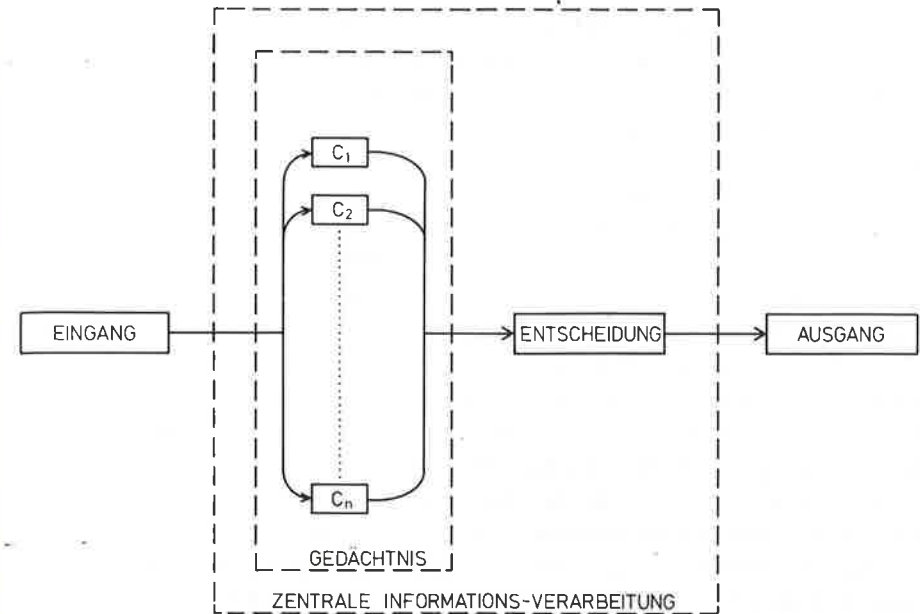


Abb. 2 Modell der Identifikation

$C_1, C_2, \dots, C_n$ sind Charakterisatoren

des menschlichen Verhaltens ansehen.

Bis jetzt haben wir offengelassen, wie das Muster mit den Charakterisatoren verglichen wird. Es sind mehrere Verfahren denkbar, und jedem entspricht eine bestimmte Anordnung der Charakterisatoren im Gedächtnisblock. Umgekehrt folgt aus dem Arrangement der Charakterisatoren ein bestimmtes Format der Resultate der Musteranalyse. Man kann die Resultate als vollständige Verzeichnisse oder Listen der Merkmale der Muster auffassen. Wir werden sie "Musternamen" nennen. Die denkbaren Anordnungsweisen der Charakterisatoren und der Namen werden wir "Modelle der Gedächtnisstruktur" nennen. Die Modelle der

Gedächtnisstruktur leiten wir im wesentlichen von den Suchstrategien ab, die wir oben als serielle und parallele Prozesse bezeichnet haben. Daher können wir sie ebenfalls in zwei Gruppen einteilen, nämlich in parallele und serielle Modelle.

Wenn die im Muster kodierte Information an alle Charakterisatoren gleichzeitig angelegt wird, handelt es sich um ein *Parallelmodell mit unbeschränkter Kapazität*. Dieses Modell wird alle Muster gleich schnell identifizieren. Sein Zeitbedarf hängt nicht von der Zahl der benutzten Charakterisatoren ab. Merkmale müssen nicht einfach sein. Allem Anschein nach sind sie verwickelte Komplexe (d.h. jedes einzelne Merkmal ist komplex, nicht nur ihre Gesamtheit) von Teilinformationen. In Untersuchungen zur Sprachwahrnehmung hat man viel Mühe aufgewandt, um die Merkmale zu ermitteln, an denen ein gesprochenes Wort und seine Laute erkannt werden - bisher vergeblich. Es hat sich herausgestellt, daß ganz unterschiedliche Geräusche als ein und derselbe Laut identifiziert werden, wenn der Kontext vorhergehender und nachfolgender Laute dies erfordert (LIBERMAN et al. 1964: 81). Die Bestimmung so komplizierter Merkmale kann hohe Anforderungen an das System stellen, z.B. was den Arbeitsaufwand beim Vergleichen des Musters mit dem Charakterisator angeht (CORCORAN 1971: 65). Das heißt, das Identifizieren wird desto länger brauchen je komplizierter das Muster ist. Wir erhalten ein Parallelmodell mit beschränkter Kapazität.

Beim Modellieren des Mustererkennens auf Rechnern verfährt man in der Regel so, daß man das dargebotene Muster sukzessiv mit allen im Programm vorgesehenen Charakterisatoren vergleichen läßt. CORCORAN (1971) hält dies Verfahren für seriell. UHR (1973) dagegen hält es für parallel, da die Identifikation erst in dem Augenblick eintritt, wenn die Resultate der Musteranalyse von allen Charakterisatoren gleichzeitig (also parallel) zur Verfügung stehen. Wir können diesen Fall als seriell-paralleles Modell bezeichnen. In diesem Modell wird offensichtlich die zur Identifikation erforderliche Zeit für alle Muster gleich sein und direkt abhängen von der Zahl der Charakterisatoren, die das Modell benutzt.

Seriell arbeitende Modelle brauchen das zu erkennende Muster nicht mit allen Charakterisatoren zu vergleichen. Es genügt, wenn nur solange weitere Charakterisatoren in Anspruch genommen werden, bis die Vorlage mit hinreichender Sicherheit identifiziert ist. Für das Resultat eines Erkennungsprozesses haben wir bereits die Bezeichnung "Mustername" eingeführt. Sei η ein Mustername, der einem Muster dann und nur dann gegeben wird, wenn es die Merkmale g , i und k besitzt. Wenn nun ein Muster vom seriellen Modell analysiert wird und die sukzessive Anwendung von Charakterisatoren zur Feststellung der Merkmale g , i und k führt, so ist klar, daß dieses Muster η ist und die Anwendung weiterer Charakterisatoren sich erübrigt.

Ein serielles Modell kann die Beziehungen zwischen Merkmalen und Namen sogar noch konsequenter ausnutzen, als das obige Beispiel erkennen läßt. Aus der Anwendung eines Charakterisators resultiert die Feststellung, ob das analysierte Muster ein bestimmtes Merkmal, sagen wir h , hat oder nicht hat. Ergibt die Analyse, daß h vorhanden ist, so muß das Muster offensichtlich den Namen haben, den alle Muster, die h enthalten, tragen. Umgekehrt, wenn h nicht vorhanden ist, so muß es denselben Namen haben, wie alle Muster, die h nicht enthalten. Im seriellen Modell kann dieser Umstand in der Weise ausgenutzt werden, daß durch die Anwendung jedes Charakterisators sukzessiv der Kreis der noch anwendbaren Charakterisatoren eingeengt wird. Ein Modell mit dieser Strategie werden wir *hierarchisches Modell* nennen; denn die in ihm ablaufende Serie von Charakterisator-Anwendungen können wir als eine Hierarchie von Tests ansehen, die den Musternamen nach und nach immer enger bestimmen.

Im hierarchischen Modell, das wir soeben beschrieben haben, bildet die Anwendungsfolge der Charakterisatoren eine Hierarchie, wie in Abb. 3 dargestellt. Die Charakterisatoren sind durch Kreise repräsentiert, die von einem Kreis ausgehenden Kanten stellen den Bereich der Charakterisatoren dar, zwischen denen nach der Anwendung des durch den fraglichen Kreis repräsentierten Charakterisators zu wählen ist, wobei die Wahl vom Resultat dieses übergeordneten Charakterisators abhängt. Die

Kreise und Kanten bilden eine baumförmige Hierarchie. Eine Anwendungsfolge von Charakterisatoren, die zur Feststellung eines Musternamens führt, entspricht einem Ast dieses Baumes. Abb. 3 veranschaulicht einen solchen Ast durch ausgefüllte Kreise und die Kanten zwischen ihnen.

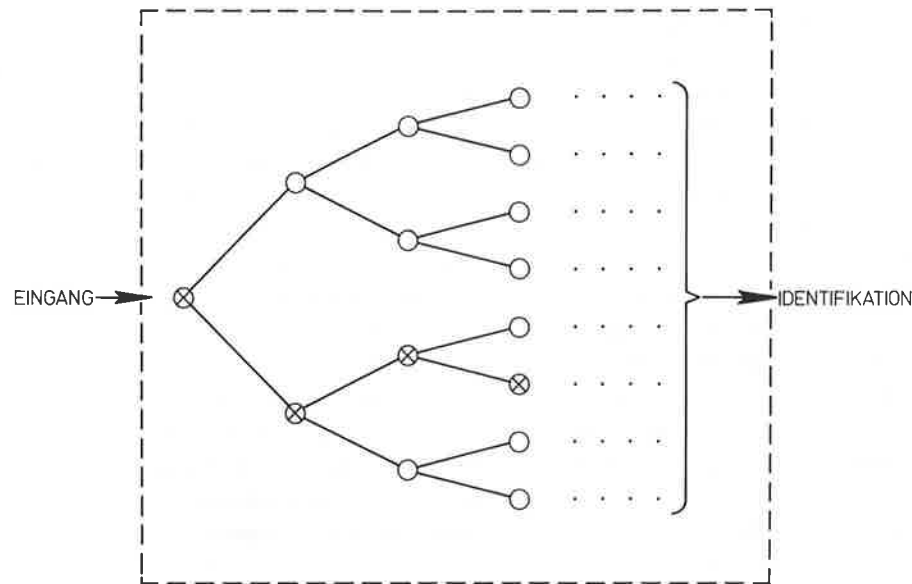


Abb. 3. Hierarchisch geordnete Charakterisatoren

Die Erkennungsgeschwindigkeit hängt im hierarchischen Modell von der Anzahl der Charakterisatoren ab, die zur Identifikation eines Musters herangezogen werden müssen. Deren Anzahl wiederum ist von der Anzahl der Namen, mit denen das Modell arbeitet, abhängig. Wie sich zeigen läßt, ist die Suchgeschwindigkeit im hierarchischen Modell direkt proportional zum Loga-

rithmus der Anzahl der Namen. Genauer: Unter der Bedingung, daß jeder Charakterisator die Zahl der noch zur Entscheidung anstehenden Namen halbiert und daß zur Festlegung jedes Namens die gleiche Anzahl von Charakterisatoren durchlaufen werden muß, gilt die logarithmische Beziehung zwischen Zeitbedarf und Umfang der Namensmenge (HICK 1952). Falls jedoch diese Bedingungen nicht erfüllt sind, werden nicht alle Muster gleich schnell identifiziert, und die Beziehung zwischen der Erkennungsgeschwindigkeit und dem Logarithmus der Namenszahl wird nicht linear sein.

2.4 Das Gedächtnis als Katalog

Die Organisation des Gedächtnisses in unseren Modellen erinnert an einen Bibliothekskatalog. Bleiben wir bei dieser Analogie, so können wir die Gesamtheit der Charakterisatoren als das Stichwortverzeichnis des Katalogs ansehen. Zwei ge-läufige Konstruktionsweisen für Inhaltsverzeichnisse von Katalogen, die hierarchische Klassifikation und das Schlüsselwortsystem, stehen dem hierarchischen bzw. dem parallelen Modell der Gedächtnisstruktur sehr nahe. Ein Beispiel für hierarchische Klassifikation ist die bekannte Dezimalklassifikation nach DEWEY. Beim Schlüsselwortsystem werden die Publikationen mit einer Anzahl von Indizes - Schlüsselwörtern - gekennzeichnet (MEADOW 1967). Beispielsweise eine Publikation aus der Entscheidungstheorie könnte mit folgenden Schlüsselwörtern indiziert werden: Entscheidung, Bayes-Ansatz, Transportprobleme.

Bei der hierarchischen Klassifikation und beim Schlüsselwortsystem erfolgt die Einordnung der Publikationen in den Katalog ähnlich wie die Mustererkennung beim hierarchischen bzw. beim parallelen Modell. Zuerst werden mit Hilfe des Katalogschlüssels die Charakteristika der Publikationen bestimmt, dann werden den Publikationen gemäß ihren Charakteristika Bezeichnungen zugeordnet - meist numerische -, die man als ihre Katalognamen ansehen kann. Die Vorstellung des Gedächtnisses als Katalog steht der Auffassung vom Gedächtnis als "cognitive

map" nahe (TOLMAN 1948, LYNCH 1960). Wenn wir nun den Speicher für faktisches und semantisches Wissen als einen Katalog ansehen könnten, so hätten wir in den Gedächtnisstrukturmodellen Theorien über seine Beschaffenheit und seinen inneren Aufbau. Es wäre überdies günstig, wenn wir von jedem der Modelle gewisse Vorhersagen über beobachtbares Verhalten der erkennenden Person ableiten könnten, um so die Gedächtnisstrukturmodelle experimentell, also empirisch zu überprüfen.

Wir haben gesehen, daß sich die Gedächtnisstrukturmodelle in der Erkennungsgeschwindigkeit unterscheiden. Man kann diese theoretischen Erkennungsgeschwindigkeiten mit denen von realen Versuchspersonen in psychologischen Experimenten vergleichen. Die Resultate solcher Vergleiche sind jedoch nicht so eindeutig, wie man vielleicht erwartet. CORCORAN (1971: 70 f.) re-sumiert eine Reihe von Experimenten, die von ähnlichen Annahmen ausgingen. Ihre Ergebnisse zeigen inkonsistent diverse Strategien der Merkmalssuche an, stützen also die verschiedenen Gedächtnisstrukturmodelle abwechselnd. Die Befunde von NEISSER (1963), NEISSER et al. (1963) und CORCORAN (1967, 1971: 88) sprechen für Parallelmodelle, die Befunde von NICKERSON und FEEHRER (1964), RABITT (1967), STERNBERG (1966, 1967) und EGETH (1966) für verschiedene serielle Modelle.

Die genannten Experimente sind jedoch nicht voll vergleichbar. Von Autor zu Autor unterscheiden sich die Modelle, die in den Experimenten überprüft werden sollten, in einer Reihe von Details. Daher kommt es, daß Ergebnisse, die ein Autor als Beleg für die serielle Strategie interpretiert, von einem anderen zugunsten des Parallelmodells gedeutet werden können. Weniger widersprüchlich sind Experimente, in denen die Versuchspersonen erst gründlich im Musteridentifizieren trainiert wurden, bevor die Erkennungsgeschwindigkeiten für einzelne Muster verglichen wurden (NEISSER u.a. 1963, CORCORAN 1967, 1971). Nach genügender Einübung werden offensichtlich alle Muster gleich schnell erkannt.

Diese Diskrepanzen zwischen den Befunden aus verschiedenen Experimenten kann man auf mehrere Weisen erklären. CORCORAN (1971: 89) nimmt an, daß beim Identifizieren unvertrauter

Muster interindividuell unterschiedliche Strategien benutzt werden. Durch die Einübung werde es den meisten Versuchspersonen jedoch möglich, zum parallelen Suchen mit unbegrenzter Kapazität überzugehen. Dieser Schluß steht im Widerspruch zu Befunden aus Experimenten, die unter dem Einfluß der Informationstheorie unternommen wurden. In ihnen wurde regelmäßig eine logarithmische Beziehung zwischen der Alternativenzahl und der Identifikationsgeschwindigkeit gefunden. Da die benutzten Reize sehr einfach waren (z.B. Aufleuchten von Signallämpchen) oder den Versuchspersonen hinreichend vertraut (wie Buchstaben oder Ziffern) und da die Anzahl der Reizalternativen immer gleich der Anzahl der Reaktionen (Namen) war, würden diese Befunde eher für das hierarchische Modell sprechen (WELFORD 1960). Übrigens schließen auch die Resultate, die CORCORAN für einen Beleg des Parallelmodells mit unbegrenzter Kapazität hält, ein hierarchisches Modell nicht aus. CORCORAN gründet seine Meinung auf die Feststellung, daß nach dem Training die Erkennungsgeschwindigkeit für alle Reizalternativen einer Gesamtheit gleich ist. Doch eine optimal geordnete Hierarchie von Charakterisatoren würde dasselbe leisten, und man könnte mit demselben Recht annehmen, daß das Training zu einer optimalen Hierarchie der Charakterisatoren geführt hat.

Die Differenzen zwischen den Ansichten der verschiedenen Autoren und die Unklarheiten in den Interpretationen experimenteller Befunde lassen sich schwerlich beheben, solange nicht das Wesen der Merkmale aufgeklärt ist. Ohne exaktere Kenntnisse über sie kann man nämlich immer annehmen, daß die experimentellen Reizvariablen bei der Wahrnehmung irgendwie umkodiert werden, so daß das, was uns als Parallelstrategie erscheint, in Wirklichkeit die serielle Strategie ist und umgekehrt (CORCORAN 1971: 63; UHR 1973: 164).

Die Analogie mit dem Bibliothekskatalog zeigt, worin der Kern des Problems liegt. Die hierarchische Klassifikation und das Schlüsselwortsystem sind zur Zusammenstellung des Katalogs gleich gut geeignet. Aber ein hierarchisch geordneter Katalog wird andere Stichwörter haben als ein nach dem Schlüsselwort-

system geordneter. Bei der Mustererkennung durch die unterschiedlichen Gedächtnisstrukturmodelle hängt die Erkennungsgeschwindigkeit vor allem von der Beschaffenheit der Charakterisatoren und damit vom Tempo der Identifikation der zugehörigen Merkmale ab. Damit man die Modelle nach der Erkennungsgeschwindigkeit vergleichen kann, müßten alle mit den gleichen Charakterisatoren arbeiten. In Wirklichkeit kann es aber so sein, daß die Beschaffenheit der Charakterisatoren von der Beschaffenheit der Merkmale bestimmt wird und daß sie ihrerseits die geeignete Organisation des Gedächtnisses bestimmt.

3. SEMANTISCHE STRUKTUREN

Die Untersuchung der Semiotik des Erkennens beginnt mit der Präzisierung der Beziehungen zwischen Gegenstand, Muster und Name. Diese Beziehungen sind allem Anschein nach analog denen in OGDEN / RICHARDS' Dreieck. Das Zeichensystem, das zu Identifikationsleistungen in Anspruch genommen wird, besteht aus einer spezifischen Art von Merkmalen, die man semantische Merkmale nennt. Den Namen hat man als den Begriff des Gegenstands bzw. als die Bedeutung des sensorischen Reizmusters aufzufassen (also nicht in dem Sinne, wie man den Ausdruck in der Semantik benutzt).

Die Theorie der semantischen Strukturen des Gedächtnisses (im folgenden kurz "Theorie der semantischen Strukturen") hat zur Erklärung von Erkennungsprozessen zwei komplementäre Auffassungen über die Gedächtnisorganisation entwickelt, nämlich eine hierarchische und eine paradigmatische Struktur. In ihnen werden die Beziehungen zwischen Gegenstand, Muster und Name unterschiedlich gesehen.

Die Theorie der semantischen Strukturen knüpft an empirische Untersuchungen zur Messung der semantischen Distanz an, wie sie in der Psycholinguistik unternommen wurden, und stellt bestimmte theoretische Forderungen an ein Distanzmaß. Gewisse Folgerungen aus dieser Theorie können empirisch überprüft werden. Die Resultate einer solchen Untersuchung ermöglichen es

zu entscheiden, ob die Struktur des Gedächtnisses hierarchisch oder paradigmatisch ist. Am Schluß des Kapitels werden einige weitere terminologische Fragen diskutiert, und es wird ein empirisches Maß der semantischen Information vorgeschlagen.

3.1 Die Semiotik des Erkennens

Im vorigen Kapitel haben wir den Begriff des Merkmals eingeführt, ohne ihn näher zu spezifizieren. Die Diskussion der Gedächtnismodelle hat jedoch gezeigt, daß zwischen der Natur der Merkmale und der Organisationsweise des an Erkennungsprozessen beteiligten Gedächtnisses ("recognition memory") eine enge Wechselbeziehung besteht. Diese Tatsache wird am klarsten beim Vergleich der Gedächtnisstruktur mit einem Bibliothekskatalog: Ein nach dem Schlüsselwortsystem angelegtes Stichwortverzeichnis kann man nicht ohne große Änderungen zur hierarchischen Klassifikation verwenden, und umgekehrt sind die Stichwörter einer hierarchischen Klassifikation nur teilweise als Schlüsselwörter brauchbar (die Beschreibungen werden umständlich).

Wüßten wir, nach welchen Prinzipien das menschliche Gedächtnis organisiert ist, so könnten wir daraus etwas über die Beschaffenheit der Merkmale folgern. Um aber Kenntnis über diese Prinzipien zu erlangen, müssen wir erst mehr über die Merkmale wissen. So drehen wir uns bei unserer Untersuchung im Teufelskreis. Wir wollen aber versuchen zu zeigen, daß es einen Ausweg gibt. Zuvor müssen wir einige Begriffe, mit denen wir arbeiten werden, präzisieren.

Bisher haben wir bei der Analyse des Erkennens vor allem eine bestimmte Seite dieses Prozesses betont, nämlich die durch Merkmale vermittelte Zuordnung eines sensorisches Reizmusters zu einem Namen. Dabei blieb der Wahrnehmungsgegenstand selbst außer Acht; aber das könnte ja eine nützliche Vereinfachung sein. Das Erkennen erscheint dann als ein einsinnig gerichteter Prozeß: Das Wahrnehmungsobjekt wird von den Sinnesorganen als sensorisches Reizmuster festgehalten. Dieses löst

dann - so die passive Theorie - die entsprechende konditionierte Antwort aus, oder es wird - so die aktive Theorie - mit Repräsentationen oder Abbildungen früherer sensorischer Information verglichen und so identifiziert. Identifikation des Musters - das ist seine Zuordnung zu einem Namen. Und der Name ist nach der passiven Deutung eine konditionierte Reaktion und nach der aktiven Deutung ein Abbild vergangener sensorischer Stimulation, dasjenige, das zum aktuellen Muster paßt.

Unter diesem Blickwinkel erscheinen Muster und Name als Anfangs- bzw. Endpunkt des Erkennungsprozesses. Zwar steht der Gegenstand noch vor dem Beginn dieser Reihe, aber der Prozeß des Erkennens fängt erst mit dem von ihm erzeugten sensorischen Reizmuster an. Zwar muß auch der Name dem anwesenden Objekt entsprechen, soll wirklich eine Erkenntnis und nicht bloß irgendeine geistige Konstruktion resultieren. Aber diese Korrespondenz von Objekt und Namen ist eine durch das Muster und den Prozeß seiner Identifikation vermittelte. Denn Information über Objekte gewinnen wir ausschließlich durch die Sinnesorgane.

Auf den ersten Blick scheint diese Auffassung vom Erkennungsprozeß richtig zu sein. Doch sie führt, wie wir schon angedeutet haben, die Erforschung des Erkennens in einen *circulus vitiosus*. Wo liegt der Fehler? Die Beschreibung bewußt auf den isolierten einzelnen kognitiven Akt zu beschränken, ist ja nicht grundsätzlich falsch. Doch setzt jeder kognitive Akt vorheriges Lernen voraus, sei dies nun als Konditionierung oder als Begriffswerwerb verstanden (Begriffe können wir vorläufig gleichsetzen mit den gespeicherten Abbildungen sensorischer Informationen gemäß ROSENBLATT's Definition der aktiven Theorie (1958)). Voll begreifen können wir das Erkennen wohl nur, wenn wir es in seiner Entwicklung, also diachron betrachten. Das bedeutet nicht, daß die Analyse des isolierten kognitiven Aktes oder einzelner seiner Phasen unmöglich oder im konkreten Fall ganz und gar ohne heuristischen Wert wäre. Wenn wir uns aber den größeren diachronen Zusammenhang der untersuchten Erscheinung nicht vergegenwärtigen, so kann die synchrone Perspektive zu verzerrten, irrigem Interpretatio-

nen führen. Vielleicht haben einige der Streitigkeiten um das Wesen des Erkennens, die wir im vorigen Kapitel besprochen haben, eben hier ihren Ursprung.

Der diachrone Standpunkt erfordert eine Umdeutung der Beziehungen zwischen Objekt, Muster und Namen. Zuerst einmal ist evident, daß die Beziehung zwischen dem Gegenstand und dem Namen direkt und nicht vermittelt ist. Sie hat sich beim vorangegangenen Lernen gebildet, als der Name als Abbildung der sensorischen Stimulation oder besser gesagt als ein gegenständlicher Begriff im Gedächtnis fixiert worden ist. Ein einzelner kognitiver Akt besteht also nicht in der Realisation der linearen Sequenz "Gegenstand - Muster - Name", sondern in der Realisation eines Tripels zweigliedriger Relationen: "Gegenstand - Name", "Muster - Name" und "Gegenstand - Muster".

Das Muster wird manchmal für eine automatische Konsequenz der Objektdarbietung gehalten. Die Befunde aus vielen Experimenten bezeugen aber etwas anderes. ZINČENKO und VERGILES (1972) belegen an mehreren Stellen ihrer Arbeit, daß die visuelle Wahrnehmung ein aktiver Prozeß ist, der von höheren Nervenzentren gesteuert wird. Nach anderen Autoren wird die Bildung eines Perzepts von einem Objekt durch eine Hypothese über dieses Objekt gesteuert. Etwas derartiges hat wohl auch LIBERMAN (1952) im Sinn, wenn er sagt, wir identifizieren beim Anhören eines Sprechers in Wirklichkeit unsere eigene unwillkürlich mitlaufende Artikulation.

Die Bildung eines sensorischen Reizmusters ist keine einfache Antwort auf einen Eindruck. Die komplizierten Mechanismen, die seine Bildung gewährleisten, werden vom Namen des dargebotenen Objekts gesteuert, natürlich nicht von ihm allein. Man kann diese Zusammenhänge am besten ausdrücken, indem man sagt, der Name vermittele die Beziehung zwischen Gegenstand und sensorischem Reizmuster. Die Situation ähnelt sehr der Situation des Bezeichnens: Ein Lautbild bezeichnet einen Gegenstand, und diese Bezeichnung wird durch einen Konzeptnamen vermittelt. Man hat in unserem Dreieck also dieselbe Anordnung von Beziehungen wie im Dreieck von OGDEN und RICHARDS, sicherlich kein Zufall. Das Erkennen ist tatsächlich ein Kommunikationsprozeß oder ge-

nauer, ein Bestandteil von Kommunikationsprozessen, in denen der Mensch die objektive Realität, die Welt der Dinge, erfährt. Man kann das Erkennen also im Rahmen der Semiotik und zwar speziell vom Standpunkt der Bedeutungstheorie, der Semantik, aus erforschen.

Die Bezeichnung "Name" ist der typischen Identifikationsaufgabe entsprechend gewählt. In ihr werden einer Versuchsperson verschiedene Objekte vorgelegt mit der Aufforderung, die Dinge zu identifizieren. Die Versuchsperson schaut sich eines der Objekte an, sagen wir den Buchstaben A, und benennt es, d.h. sie sagt "A". Die Bedeutung, die wir hier dem Ausdruck "Name" geben, unterscheidet sich ziemlich von der ursprünglichen Bedeutung des Wortes und vielleicht noch mehr von seiner Verwendung in der Semantik. In unserem Dreieck "Gegenstand - Muster - Name" bezeichnet der Name ein Konzept des Gegenstands, allerdings eine etwas ungewöhnliche Art von Konzept. In der Regel wird Konzept als Allgemeinbegriff verstanden. Hier aber handelt es sich um das Konzept eines individuellen Objekts. Um dies auch terminologisch zum Ausdruck zu bringen, sprechen wir nicht vom "Konzept", sondern vom "Namen". In Bezug aufs Muster ist der Name eine Bedeutung, nämlich die vom Muster ausgelöste Bezeichnungsreaktion. Durch diese Zuordnung des Musters zum Namen wird die Bezeichnung des Gegenstands fixiert.

Die Semiotik erforscht Zeichensysteme. In unserem Fall sind die Zeichen vom Typus der Merkmale. Üblicherweise faßt man Merkmale als sensorische Korrelate von Dingeigenschaften auf, aber diese Auffassung ist allzu weit. Es sind ja nicht alle sinnlich widergespiegelten Eigenschaften für die Identifikation des Dinges gleich bedeutsam. Im 2. Kapitel (S.115) haben wir den Namen als Merkmalsliste bezeichnet. Eine solche Liste kann wohl kaum ein vollständiges Verzeichnis aller sensorisch abbildbaren Dingeigenschaften sein. Sie muß aber ein vollständiges Verzeichnis derjenigen Merkmale sein, die für das gegebene Ding typisch sind und seine Unterscheidung von anderen Dingen ermöglichen. Im folgenden werden wir die Namen bildenden Merkmale "semantische Merkmale" nennen.

Man beachte, daß die hier eingeführte terminologische Präzisierung und die semiotische Ansicht vom Erkennen nicht mit der im 2. Kapitel (S.115) dargelegten einfacheren Auffassung im Widerspruch stehen. Die semiotische Analyse des Erkennens hat zur Präzisierung der dort verwendeten Ausdrücke geführt, und sie ermöglicht es, die untersuchten Probleme in neuen Zusammenhängen zu sehen. Die Beteiligung des Namens an der Musterbildung ist allerdings ein weniger klarer Aspekt der semiotischen Betrachtungsweise geblieben. Die Sinnesorgane registrieren physikalische und chemische Reize, denen dann Merkmale des Musters entsprechen. Zum Identifizieren sind jedoch nur Merkmale im engeren Sinn, semantische Merkmale, wesentlich. Die Funktion des Namens bei der Vermittlung der Beziehungen zwischen Objekt und Muster wird also in der Selektion und Umkodierung von Merkmalen in semantische Merkmale bestehen.

3.2 Semantische Strukturen

Wenn die Namen Konzepte der zu identifizierenden Objekte sind, so müßte es möglich sein, Experimente zu entwerfen, die wenigstens Teileinsichten in ihre Anordnung im individuellen Gedächtnis erbringen. Mit "Anordnung" meinen wir hier nicht eine anatomische oder geometrische Anordnung im realen Raum, sondern eine funktionelle Anordnung. Deren Topologie wollen wir untersuchen. Eine Theorie der semantischen Strukturen ist eine Theorie der funktionellen Organisation der Namen. Im folgenden Kapitel werden Experimente beschrieben, die eine solche Theorie stützen sollen. Wenn wir sagen, die Namen seien Konzepte, so kennzeichnen wir damit ihre Funktion, ihre Rolle im Erkennungsprozeß. Hinsichtlich der Frage, woraus und wie sie gebildet werden, vertreten wir die Ansicht, daß die Namen semantische Merkmalslisten sind. Bezeichnen wir als J die Gesamtmenge oder eine Teilmenge der im individuellen Gedächtnis gespeicherten Namen und als Φ die Menge aller semantischen Merkmale, die auf den Listen vorkommen, welche die in der Menge J enthaltenen Namen bilden. Wenn zwei dargebotene Objekte

eine gemeinsame Eigenschaft haben, z.B. wenn beide rot sind, so werden die ausgelösten sensorischen Reizmuster ein bestimmtes semantisches Merkmal gemeinsam haben. Unterschiedliche Namen aus der Menge J werden doch einige semantische Merkmale gemeinsam haben. Wir können daher auf dem kartesischen Quadrat der Menge J eine Ähnlichkeitsfunktion einführen, welche die Ähnlichkeit zweier Namen als die Anzahl ihrer gemeinsamen semantischen Merkmale definiert.

Zu der Ähnlichkeitsfunktion kann man eine Umkehrfunktion definieren, d.h. ein Maß für die Unterschiedlichkeit oder Entfernung je zweier Namen. Wir werden fordern, daß diese Funktion eine Metrik ist, und bezeichnen sie mit ρ . Das bedeutet, daß ρ eine symmetrische und reflexive Funktion ist, die folgende Eigenschaften hat: a) Für zwei beliebige Namen x und y aus der Menge J gilt $\rho(x,y) = 0$ dann und nur dann, wenn $x = y$ ist; b) Für jedes Tripel x, y, z aus der Menge J gilt die Dreiecksungleichung:

$$\rho(x,y) \leq \rho(x,z) + \rho(y,z).$$

(Ein Existenzbeweis für dieses Distanzmaß ρ wird im 6. Kapitel bei der Erörterung der Theorie semantischer Strukturen geben.)

Die Funktion ρ werden wir als *Maß der semantischen Distanz* bezeichnen. Die Messung der semantischen Distanz hat in der psychologischen Literatur schon eine gewisse Tradition. Wir begegnen ihr des öfteren in der Psycholinguistik; doch OSGOOD et al. (1957) und ihre Nachfolger haben sie auch zur Untersuchung von Einstellungen zu Kunstwerken, Fachzeitschriften und politischen Fragen mit Erfolg eingesetzt (vgl. z.B. KÜLLER 1972). Wir bezeichnen mit "Metrisierung" die theoretischen und mit "Messen" die empirischen Operationen. Gegenwärtig und besonders im 6. Kapitel betrachten wir die Metrisierung der semantischen Distanz als eine theoretische Operation. Der empirischen Seite des Messens ist ein großer Teil des 4. Kapitels gewidmet. Dort handelt es sich aber um Messung im schwachen Sinne, d.h. "Distanz" kann nur im uneigentlichen

und nicht im geometrischen Sinne verwendet werden.

Eine *semantische Struktur* ist gegeben durch das Tripel $\langle J, \Phi, \rho \rangle$, d.h. die Menge der Namen, die Menge der semantischen Merkmale und das semantische Distanzmaß über dem kartesischen Produkt der Namenmenge. Wir zeigen nun, daß man zwei Typen von semantischen Strukturen, d.h. zwei Arten der funktionellen Gedächtnisorganisation der Namen, unterscheiden kann, je nachdem, welche metrischen Beziehungen in einer Teilmenge von Namen aus J mit einem gemeinsamen Merkmal aus Φ gelten.

Seien δ eine beliebige nichtnegative reelle Zahl und i ein Element aus J . Als δ -Umgebung von i bezeichnen wir eine solche Namenmenge J' , $J' \leq J$, daß für jedes Element j aus J' gilt: $\rho(i,j) \leq \delta$.

Wenn nun ein hierarchisches Gedächtnismodell vorliegt, so sind seine Charakterisatoren wie in einem Stammbaum geordnet und erzeugen folglich eine hierarchische Ordnung der semantischen Merkmale. Da die semantische Distanz eines Namenpaares eine Funktion der Anzahl gemeinsamer Merkmale ist, hat eine δ -Umgebung eines Namens im hierarchischen Modell besondere Eigenschaften. Nehmen wir zwei Namen mit einem gemeinsamen Merkmal, das von einem Charakterisator auf einem bestimmten Niveau der Hierarchie determiniert wird. Aus der hierarchischen Anordnung der Charakterisatoren folgt mit Notwendigkeit, daß die beiden Namen auch alle diejenigen Merkmale gemeinsam haben, deren zugehörige Charakterisatoren auf der Kante zwischen dem Charakterisator des gegebenen Merkmals und der "Quelle" (Stammbaumwurzel, also dem obersten Charakterisator) liegen. Daraus ergibt sich nun, daß die einer δ -Umgebung entsprechende Namenmenge als Ganze einen bestimmten Satz von Merkmalen gemeinsam hat, der von der Größe der Zahl δ abhängt. Dies gilt aber nur im hierarchischen Modell.

Die Gesamtheit der Namen kann derart in Teilmengen zerlegt werden, daß die Elemente einer Teilmenge jeweils durch ein gemeinsames Merkmal gekennzeichnet sind.

Eine *semantische Struktur* ist *hierarchisch* genau dann, wenn zu jedem ihrer semantischen Merkmale eine nichtnegative reelle Zahl δ existiert, so, daß die δ -Umgebung jedes Namens

mit dem gegebenen Merkmal alle und nur die Elemente der durch dieses gemeinsame Merkmal gekennzeichneten Teilmenge enthält.

Die hierarchischen Strukturen haben zwei bedeutsame Eigenschaften. Erstens ist jede hierarchische Struktur stammbaumförmig. Zweitens ist das in einer hierarchischen Struktur geltende semantische Distanzmaß ρ eine Ultrametrik. Eine Ultrametrik ist ein Sonderfall einer Metrik, in dem anstelle der Dreiecksungleichung eine stärkere Beziehung gilt. Bezeichnen wir mit i, j, k drei beliebige Namen und mit $\rho(x,y)$ den Wert des semantischen Distanzmaßes eines Namenpaares. Die ultrametrische Ungleichung, die im Falle einer Ultrametrik die Dreiecksungleichung ersetzt, legt fest, daß der Abstand zweier Namen höchstens so groß ist wie der größere ihrer Abstände zu einem dritten Namen, nämlich:

$$\rho(i,j) \leq \max [\rho(i,k), \rho(j,k)]$$

Die hierarchische Struktur ist die erste der beiden funktionellen Organisationsweisen von Namen im Gedächtnis, welche wir im Rahmen der Theorie der semantischen Strukturen unterscheiden können. Die zweite Organisationsweise bezeichnen wir als paradigmatisch.

Eine *semantische Struktur* heißt *paradigmatisch*, wenn sie nicht die Definition einer hierarchischen Struktur erfüllt.

Eine paradigmatische Struktur ist nicht stammbaumförmig. In ihr ist das semantische Distanzmaß eine gewöhnliche Metrik, für die die Dreiecksungleichung gilt.

Die semantische Erweiterung der Gedächtnisstrukturmodelle ermöglicht es - wenigstens in der oben eingeführten Auffassung - das hierarchische Modell von den anderen Modellen zu unterscheiden. Diese lassen sich unter dem Titel "paradigmatische Strukturen" zusammenfassen. Das folgt doch aus dem Zusammenhang des semantischen Distanzmaßes mit der Beschaffenheit der erörterten Modelle. Das hierarchische Modell unterscheidet sich von allen anderen gerade dadurch, daß es als einziges beim Erkennungsprozeß die innere Struktur der Namensgesamtheit konsequent ausnützt. Die anderen Modelle - gleich, ob

sie die serielle oder die parallele Strategie verwenden - nutzen die gegenseitigen Beziehungen zwischen den Namen nicht aus und bilden so, vom Standpunkt der Theorie semantischer Strukturen aus gesehen, in ihrer Gesamtheit ein zum hierarchischen Modell komplementäres Modell. Die semantischen Strukturen modellieren also einen etwas anderen Aspekt des Erkennungsprozesses als die Modelle des parallelen und des seriellen Suchens. Dieser Aspekt ist jedoch nicht weniger wichtig, überdies hängt er eng mit den Problemen der Gedächtnistheorie zusammen.

Stellen wir uns ein einfaches Gedächtnisexperiment vor. Wir legen der Versuchsperson einige Dinge vor und bitten sie, sich diese zu merken. Zunächst einmal muß die Versuchsperson die dargebotenen Objekte erkennen, also die Reizmuster ihren Namen zuordnen. Diese Namen bzw. die sie bildenden semantischen Merkmalslisten können dann im Kurzzeitgedächtnis gespeichert werden. Sie können einige Merkmale gemeinsam haben. Falls die Struktur des "recognition memory" paradigmatisch ist, müssen die Merkmalslisten als Ganze erhalten bleiben, andernfalls wäre die Reproduktion unmöglich. Wenn jedoch die Gedächtnisstruktur hierarchisch ist, sind der Namensgesamtheit eine gewisse Anzahl Merkmale (mindestens das oberste in der benutzten Hierarchie) gemeinsam (siehe Kap. 6). Diese tragen nichts zur Unterscheidung bei und können daher bei der Speicherung der Namen eingespart werden.

In einer semantischen Struktur kann man zu einer beliebigen Menge von Namen, deren gegenseitige Abstände bekannt sind, eine δ -Umgebung konstruieren, indem man als δ den größten Abstand wählt. In einer hierarchischen Struktur gilt dann außerdem: Je kleiner δ ist, desto mehr gemeinsame Merkmale hat die Namensmenge. In einer paradigmatischen Struktur besteht nur jeweils für die einzelnen Paare ein Zusammenhang zwischen ihrem inneren Abstand und der Anzahl ihrer gemeinsamen Merkmale. Wenn wir also beispielsweise feststellen, daß das Einprägen einer größeren Zahl von Objekten umso leichter ist, je kleiner die semantischen Abstände dieser Objekte sind, so können wir schließen, daß die Struktur des beteiligten "recognition memory" hierarchisch ist.

3.3 Einige Anmerkungen

In der Theorie der semantischen Strukturen werden die hierarchische und die paradigmatische Struktur als stabile Anordnungen der Namen und der semantischen Merkmale dargestellt. Wie wir jedoch oben gesagt haben, ist eine semantische Struktur die Struktur eines sich entwickelnden flexiblen Systems. Um diese beiden Ansichten in Einklang zu bringen, können wir die semantische Struktur als eine allgemeine Anweisung zur fortlaufend erneuerten funktionsgerechten Anpassung des sich entwickelnden Systems der Erkennungsprozesse auffassen. Von diesem Standpunkt aus ist eine semantische Struktur eine generelle Anweisung zur Gewährleistung von Erkennungsprozessen oder, wie man vielleicht sagen kann, eine generelle Erkennungsmethode.

Als wir in diesem Kapitel den Ausdruck "Name" präzisierten, unterschieden wir ihn vom Konzept, das insofern allgemeiner ist, als ein Name das Konzept von einem individuellen Gegenstand ist. Doch nicht immer richtet sich das Erkennen auf ein individuelles Objekt. In vielen natürlichen Situationen reicht es, wenn ein Objekt als Element einer Klasse bestimmt wird. Um ein Beispiel von BRUNER et al. (1967) zu nehmen: Ein Amateur-Pilzsammler ist zufrieden, wenn er eßbare Pilze von ungenießbaren unterscheiden kann. Die einzelnen Pilzarten braucht er kaum zu kennen und Pilzindividuen zu identifizieren, ist wohl überflüssig.

In der Theorie der semantischen Strukturen hat eine Klasse keinen Namen, sondern sie kann als die Liste derjenigen semantischen Merkmale bestimmt werden, die den Elementen der Klasse gemeinsam sind. So ist es möglich, ein semantisches Merkmal als ein Konzept einer Eigenschaft oder als ein Konzept einer Klasse von Objekten, die alle diese Eigenschaft haben, zu sehen. Man kann auch zeigen, daß es nicht nötig ist, konsequent auf der Forderung zu beharren, daß ein Name das Konzept von einem individuellen Gegenstand sei. Jedoch muß in jeder konkreten semantischen Struktur konsequent festgelegt sein, welche von den sie bildenden Konzepten Merkmale und welche Namen

sind; dies aus analogen Gründen, wie sie in der Logik zur Typentheorie geführt haben.

In der Theorie der kognitiven Prozesse wie in der Gedächtnistheorie hat vor einiger Zeit die SHANNONSche Informationstheorie eine bedeutende Rolle gespielt. Gegen ihre Anwendung in der Psychologie wurde unter anderem oft der Einwand erhoben, daß der Mensch nicht mit probabilistisch definierten Signalen, sondern mit Bedeutungen arbeite. Daher wird in der Psychologie zunehmend der semantischen Information mehr Bedeutung beigemessen als der von SHANNON definierten probabilistischen. Die Theorie der semantischen Strukturen ermöglicht die Einführung eines empirischen Maßes der semantischen Information. Man kann den Erkennungsprozeß als einen Selektionsprozeß sehen. Aus der Menge von Möglichkeiten, die wir hier mit der Menge J der Namen identifizieren wollen, wird durch die Zuordnung des Reizmusters zu einem Namen eine einzige Möglichkeit selektiert. So wird die zu Beginn des Erkennungsprozesses bestehende Unbestimmtheit auf Null reduziert, d.h. in Gewißheit verwandelt. Die Selektion einer Teilmenge aus der Menge der Namen kann ebenfalls als Unbestimmtheitsreduktion aufgefaßt werden. Hier wird jedoch die resultierende Unbestimmtheit größer als Null bleiben.

Der größte Abstand innerhalb einer Menge von Elementen heißt Durchmesser der Menge. Bezeichnen wir mit J die Menge der Namen, mit J_1 eine beliebige ihrer Teilmengen und $d(J)$ bzw. $d(J_1)$ die Durchmesser der Mengen. Durch die Selektion einer Teilmenge wird die mit der Gesamtmenge verbundene semantische Unbestimmtheit reduziert. Das Ausmaß der Unbestimmtheitsreduktion, die bei der Selektion von J_1 aus J erfolgt, bezeichnen wir mit $v(J_1/J)$. Diese Größe kann durch die folgende Beziehung ausgedrückt werden:

$$v(J_1/J) = \frac{d(J) - d(J_1)}{d(J)} .$$

Die Funktion v nimmt Werte im Intervall $<0,1>$ an; $v = 1$ ergibt sich, wenn das Resultat der Selektion genau ein Name ist, $v = 0$, wenn $J_1 = J$ ist. Die Funktion v ist kein direktes Maß

der semantischen Information; doch man kann sie sicherlich als Schätzung des von einer Teilmenge getragenen semantischen Informationsbetrags oder zum Vergleich der von mehreren Teilmengen einer semantischen Struktur getragenen Informationsbeträge verwenden.

Wir haben in diesem Kapitel verschiedentlich über Semiotik und Semantik gesprochen. Weder die Semiotik noch die Semantik ist eine einheitliche Wissenschaftsdisziplin. Eine Erörterung der verschiedenen semiotischen und semantischen Theorien ist jedoch nicht Aufgabe dieser Studie. Der Autor ist allerdings der Meinung, daß sich die Beziehungen zwischen Objekt, sensorischem Reizmuster und Objektkonzept, von denen die Identifikationstheorie ausgeht, nur unter einem semiotischen bzw. semantischen Blickwinkel begreifen lassen. Aus dem Versuch, diese Beziehungen, soweit sie im Erkennungsprozeß manifest werden, theoretisch zu erfassen, ist eine Theorie der semantischen Strukturen entstanden, die zwei alternative Auffassungen von diesen Beziehungen umfaßt (die hierarchischen und die paradigmatischen Strukturen). Die Semiotik der Wahrnehmung, wie sie hier dargestellt wurde, stützt sich auf die Theorie der Signifikation.

4. MESSUNGEN UND EXPERIMENTE

Die Theorie der semantischen Strukturen ist eine Theorie über die Gedächtnisorganisation von eingprägtem Material. Wie wir im vorigen Kapitel gezeigt haben, kann man aus ihr einige Vorhersagen über Gedächtnisleistungen ableiten. In diesem Kapitel werden wir Experimente beschreiben, die diese Vorhersagen überprüfen sollen. Zuvor müssen wir jedoch noch einige Fragen klären. Als erste stellt sich die Frage nach der Allgemeingültigkeit der Theorie der semantischen Strukturen. Bezieht sich die Theorie der semantischen Strukturen auf das Behalten aller Arten von Reizmaterial oder nur auf das Behalten bestimmter Arten von Material, z.B. verbalen Materials? Diese Frage ist für die Theorie semantischer Strukturen entscheidend.

Wenn wir nämlich die innere Organisation des Gedächtnisses ermitteln wollen, so müssen wir die gegenseitigen Beziehungen einer großen Menge von eingprägten Items untersuchen. Denn wenn nur eine geringe Anzahl von Items untersucht wird, ist die Gefahr nicht auszuschließen, daß die erhaltene Struktur nur ein Artefakt ist, das im wesentlichen von der Itemauswahl abhängt. Wenn wir nun verbales Material benutzen, können wir verhältnismäßig leicht mit großen Itemmengen arbeiten. Bei nichtverbalem Material wäre das ziemlich schwierig, wenn nicht direkt unmöglich.

Ein weiteres Problem ist mit der Messung der semantischen Distanz verbunden. Semantische Distanzen zu messen, ist für die Überprüfung der Theorie der semantischen Strukturen unumgänglich. Es ist also nötig, eine geeignete Meßmethode zu finden und darüber hinaus zu zeigen, daß das semantische Distanzmaß die Eigenschaften hat, die wir ihm zuschreiben, nämlich daß es eine Metrik ist und daß es auf der Anzahl der einem Namenpaar gemeinsamen semantischen Merkmale basiert.

Das erste Problem, also die Frage nach der Allgemeingültigkeit der Theorie semantischen Strukturen, können wir wohl nicht endgültig entscheiden.

Die Lösung der zweiten Frage, nämlich die Wahl einer geeigneten Methode zur Messung der semantischen Distanz, stellt uns nicht so anspruchsvolle Probleme. FILLERBAUM (1971: 290) nennt zwei grundlegende Richtungen in diesem Gebiet, die er einerseits durch OSGOOD und seine Anhänger, andererseits durch G.A. MILLER (MILLER & McNEILL 1969) repräsentiert sieht. Nach der Ansicht CLIFF's (1973) ist das Semantische Differential weniger exakt als andere Skalierungsmethoden. Außerdem ist OSGOOD's Methode auf die Untersuchung der Konnotation gerichtet. Für die Theorie der semantischen Strukturen ist solch eine Einschränkung ungünstig. Deswegen verwenden wir die von MILLER vorgeschlagene Methode. Diese wird unten ausführlicher beschrieben. Ein Nachweis, daß dieses Maß von der Anzahl der den Namen paarweise gemeinsamen semantischen Merkmale abhängt, ist an die Messung von Gedächtnisleistungen gebunden. Wenn es uns gelingt zu zeigen, daß zwei Namen desto besser behalten werden,

je kleiner ihre semantische Distanz ist, so bedeutet das offenbar, daß beim Einprägen benachbarter Namen eine Ersparnis eintritt, die sich am einfachsten durch Merkmalsgemeinsamkeit erklären läßt. Die auf beiden Listen vorkommenden Merkmale brauchen nicht bei jedem Namen, sondern nur einmal eingepreßt zu werden. Diese Annahme impliziert dann die Umkehrung: Die semantische Entfernung von Namen ist desto geringer, je mehr semantische Merkmale sie gemeinsam haben. Überprüfen könnte man die Ersparnishypothese z.B. mit einem Experiment, das die Methode des Paarassoziationslernens einsetzt.

Das Hauptziel der in diesem Kapitel darzustellenden Experimente ist jedoch ein Versuch herauszufinden, ob die semantische Struktur des menschlichen Gedächtnisses hierarchisch oder paradigmatisch ist. Im vorigen Kapitel haben wir bereits skizziert, wie die Experimente angelegt sein müßten, mit denen das möglich wäre. Beim Behalten einer größeren Menge von Namen - sagen wir, k größer als zwei - entsteht, sofern eine hierarchische Struktur beteiligt ist, eine desto größere Ersparnis, je geringer die gegenseitigen semantischen Distanzen der Namen in dieser k -Gruppe sind. In einer paradigmatischen Struktur gilt diese Beziehung nicht. Wie schon erwähnt, benutzen wir in den Experimenten, in denen wir diese Beziehungen untersuchen, verbales Material, und zwar tschechische Substantive.

4.1 Messung der semantischen Distanz

Die in dieser Arbeit verwendete Methode zur Messung der semantischen Ähnlichkeit von Substantiven ist von G.A. MILLER (1969) formuliert worden. Die Messung beruht auf der Sortierung der Substantive nach Bedeutungsähnlichkeit. MILLER verwendete 48 englische Substantive, von denen jedes zusammen mit einer die Bedeutung spezifizierenden Kurzdefinition und einem den Gebrauch illustrierenden einfachen Satz auf ein Kärtchen gedruckt war. Die 48 Kärtchen legte er seinen Versuchspersonen mit der Anweisung vor, sie nach Bedeutungsähnlichkeit zu

gruppieren. MILLER setzte keine Beschränkungen hinsichtlich der Zahl und Größe der Grüppchen fest und gab den Vpn auch keine Erklärung darüber, was er mit Bedeutungsähnlichkeit meinte.

Die von den einzelnen Vpn gelieferten Gruppierungen trug MILLER in eine Matrix der Größe 48×48 ein. Die Zelle in der i -ten Spalte und der j -ten Zeile der Matrix repräsentierte das Namenpaar i und j . In der Zelle i, j wurde eine 1 notiert, wenn beide Namen in dasselbe Grüppchen sortiert worden waren. Der Wert 0 wurde notiert, wenn die Vpn die Substantive i und j unterschiedlichen Grüppchen zugeordnet hatte. MILLER erhielt zunächst für jede individuelle Sortierung eine Matrix dieser Art und addierte die Matrizen dann. Es entstand eine Matrix, deren i, j -te Zelle angab, in wievielen der individuellen Sortierungen das Paar i, j im gleichen Grüppchen vorkam.

Diese Matrix ist symmetrisch; denn man kann nicht i mit j zusammenlegen, ohne daß zugleich j mit i zusammenfiele. Die Hauptdiagonale kann weggelassen werden, da alle Werte in ihr gleich der Vpn-Zahl sind. Die in den Zellen eingetragenen Häufigkeiten können mit der semantischen Ähnlichkeit oder Nähe der Substantive in Beziehung gesetzt werden. Setzen wir nämlich N gleich der Anzahl der Vpn und N_{ij} gleich dem Wert in der Zelle i, j , so können wir die semantische Distanz zweier Namen i, j durch die folgende Gleichung ausdrücken:

$$D_{ij} = N - N_{ij}.$$

Wir wollen zeigen, daß D_{ij} eine Metrik ist. Zunächst fällt offensichtlich bei jeder Sortierung jeder Name mit sich selbst zusammen. Daher ist $D_{ii} = 0$. Wegen der Symmetrie ist $D_{ij} = D_{ji}$. In jeder individuellen Matrix ist auch die Dreiecksungleichung erfüllt: $D_{ik} \leq D_{ij} + D_{jk}$; denn man kann nicht i und j sowie gleichzeitig j und k zusammenlegen, ohne daß auch i und k zusammenfallen. Bei der Addition der Matrizen bleiben diese Beziehungen erhalten. Daher ist D auch für die Summenmatrix eine Metrik.

Die so gewonnenen Distanzdaten verarbeitete MILLER mit Hilfe von JOHNSON's (1967) hierarchischer Gruppierungsmethode. Diese Methode stellt eine Distanzmatrix, in der die ultrametrische Ungleichung erfüllt ist, als einen Stammbaumgraphen dar. Dabei wird nur eine Partialordnung der Distanzen vorausgesetzt.

MILLER stellte sich das Ziel, Einblick in die Beschaffenheit der semantischen Merkmale von sprachlichen Konzepten zu gewinnen. Von ihm stammt die Unterscheidung hierarchischer und paradigmatischer Strukturen. Er bezog sie allerdings nur auf die skalierten Distanzen von Substantiven. Doch durch Skalierung allein kann zwischen den beiden Strukturtypen nicht empirisch entschieden werden. Selbst wenn gewährleistet wäre, daß das semantische Distanzmaß eine Ultrametrik ist, so wäre das im Gegensatz zu MILLER's Vermutung doch keine Garantie dafür, daß eine hierarchische Struktur vorliegt.

Da die Ultrametrik in der Theorie der semantischen Strukturen einen wichtigen Platz einnimmt, führen wir noch einige Details von JOHNSON's Methode an. Die hierarchische Gruppierung wird nach einem Algorithmus durchgeführt, mit dessen Hilfe eine Abbildung der Distanzmatrix erzeugt wird. Nehmen wir die Distanzmatrix, die der in MILLER's Experiment erhaltenen semantischen Ähnlichkeitsmatrix entspricht. Da die Matrix symmetrisch ist und die Werte in der Hauptdiagonalen weggelassen werden, betrachten wir die Werte oberhalb der Hauptdiagonalen näher. Wir suchen hier den kleinsten Wert. Die ihm entsprechenden Wörter lassen wir als eine Einheit in eine neue Matrix eingehen, welche eine Spalte und eine Zeile weniger hat. Diesen Vorgang wiederholen wir solange, bis alle Spalten und Zeilen der ursprünglichen Matrix aufgebraucht sind. Die gebildeten Verklumpungen von Wörtern werden als Knoten in einem Stammbaumgraphen dargestellt.

Falls die Daten in der Matrix die ultrametrische Ungleichung erfüllen, wird die beschriebene Prozedur glatt verlaufen; denn für beliebige i, j, k und $D_{ij} \leq D_{ik}$, $D_{ij} \leq D_{jk}$ wird $D_{ik} = D_{jk}$ sein. Wenn aber die Daten durch Zufallsstörungen verzerrt sind, entsteht das Problem, wie man die Distanz $d((ij)k)$

zwischen dem Klumpen (ij) und dem Wort k in den Fällen bestimmen soll, wenn $D_{ik} \neq D_{jk}$ ist.¹ Für diese Fälle schlägt JOHNSON vor, das Problem doppelt zu lösen. Für die erste Lösung verwendet man immer den kleineren der beiden nichtidentischen Werte, für die zweite immer den größeren. Die beiden Lösungen ergeben zwei hierarchische Graphen. Wenn sich die beiden Graphen nicht allzu sehr voneinander unterscheiden, können wir annehmen, daß in den analysierten Daten die ultrametrische Ungleichung annähernd gilt. Andernfalls ist entweder die analysierte Struktur nicht hierarchisch, oder die Daten sind durch Zufallsstörungen so verzerrt, daß man keine exakte Analyse durchführen kann (MILLER 1969: 181).

Die beiden beschriebenen Verfahren bezeichnet JOHNSON als Maximum- und als Minimum-Methode. Für die Maximum-Methode benutzt er auch die Bezeichnungen "diameter method" und "clustering by complete linkage", für die Minimum-Methode die Bezeichnungen "connectedness method" und "clustering by single linkage". Die Distanz eines Elements z von einem Cluster (x, y) ist bei der Maximum-Methode definiert als

$$d((x, y), z) = \max (d(x, z), d(y, z))$$

und bei der Minimum-Methode als

$$d((x, y), z) = \min (d(x, z), d(y, z)).$$

Bei der Maximum-Methode wird also der Durchmesser der Vereinigungsmenge von x , y und z bestimmt, und bei der Minimum-Methode wird der Abstand von z zum nächstgelegenen Element des Clusters (x, y) ermittelt. In beiden Verfahren wird das Element z so ausgewählt, daß der ermittelte Kriteriumswert - Clusterdurchmesser bzw. Abstand zum nächstgelegenen Element - minimal wird.

4.2 Die Experimente

Die Experimente, die wir hier beschreiben wollen, wurden im Laboratorium des Psychologischen Instituts der Tschechoslowakischen Akademie der Wissenschaften durchgeführt. Die Ver-

¹ Mit d_{ij} werden Werte der theoretischen Distanzfunktion d bezeichnet, mit D_{ij} dagegen die empirischen Zahlen $N-N_{ij}$.

suchspersonen waren Studenten und Studentinnen. Sie wurden für ihre Teilnahme an den Experimenten bezahlt.

Das Untersuchungsziel war, einen Zusammenhang zwischen den semantischen Distanzen von Substantiven und den Behaltensleistungen beim Erlernen dieser Substantive nachzuweisen. Die Messung der semantischen Distanzen wurde zweimal durchgeführt und die Behaltensleistungen wurden mit drei verschiedenen Verfahren gemessen: mit dem Verfahren der freien Reproduktion, mit dem Verfahren des Wiedererkennens und mit der Paarassoziationsmethode.

Die Versuche wurden gruppenweise mit je ein bis fünf Personen durchgeführt. Alle Teilnehmer hatten gutes Licht und genügend Platz zum Umgang mit den Vorlagen. Zu Beginn wurde den Vpn gesagt, daß sie an mehreren Versuchen teilnehmen würden, von denen die meisten ihre Gedächtnisfähigkeit beanspruchten. Es wurde dabei betont, daß die Daten nicht im Hinblick auf ihre individuellen Leistungen, sondern im Hinblick auf Eigenschaften des vorgelegten Materials interpretiert werden sollten.

Nach dieser Einführung hatten die Vpn zwei Tests zu bearbeiten, bei denen sie mit dem ungewohnten Milieu und der Versuchssituation vertraut werden sollten: die Progressiven Matrizen von RAVEN sowie einen sprachlichen Gedächtnistest (eine Modifikation eines Verfahrens, das von den Tschechoslowakischen Staatsbahnen bei Neueinstellungen verwendet wird). Darauf folgten die drei Gedächtnisversuche und die beiden semantischen Distanzmessungen. Die eine Distanzmessung ging den Gedächtnisversuchen unmittelbar voran, die andere folgte ihnen unmittelbar.

Insgesamt 49 Personen nahmen an den Versuchen teil. Wir erhielten im ganzen je 49 Protokolle von den beiden Messungen der semantischen Distanz und dem Paarassoziationslernen sowie je 42 Protokolle aus den Versuchen mit freier Reproduktion und Wiedererkennen.

4.2.1 Durchführung der semantischen Distanzmessung

Die semantischen Distanzen wurden an einer Stichprobe von 52 Substantiven gemessen. 48 von diesen Wörtern waren tschechische Übersetzungen der von MILLER (1969) benutzten Wörter. 4 weitere Wörter waren in einem Vorversuch so ausgewählt worden, daß sie mit einer zuvor festgelegten Gruppe von 6 Wörtern des schon vorhandenen Materials ein dicht zusammenhängendes Cluster bildeten. (Das war nötig, um die Versuchspläne für das freie Reproduzieren und das Wiedererkennen zu erfüllen.)

Jedes der 52 tschechischen Substantive (darunter waren zwei Substantive mit einem Adjektiv) wurde mit Maschine auf ein festes Pappkärtchen im Format 7 x 10.5 cm geschrieben. Das Wort wurde etwas oberhalb der Mitte des Kärtchens geschrieben, und zwar in Großbuchstaben und gesperrt. Darunter wurde in Normalschrift eine Definition getippt. Im Unterschied zu MILLER wurden keine Sätze hinzugefügt, die den Gebrauch der Wörter illustrieren (siehe Tab. 1).

Tab. 1: Die in den Experimenten verwendeten 52 Substantive und ihre Definitionen

1. ANKER, ein Eisenstück bestimmter Form, an Kette oder Seil befestigt; soll ein Schiff am Ort festhalten
2. BLEICHMITTEL, eine Chemikalie, die zum Bleichen verwendet wird
3. KOCH, ein Mensch, der Speisen zubereitet
4. ARZT, eine Person mit der Berechtigung, Krankheiten zu behandeln
5. AUSPUFFGAS, entweichender Rückstand von verbranntem Benzin
6. FISCH, Wasserlebewesen mit Kiemen
7. LEIM, Stoff, der zum Kleben von Sachen benutzt wird
8. HECKE, ein schmaler kultivierter Streifen von Büschen, der als Zaun dient
9. BÜGELEISEN, Gerät zum Bügeln von Kleidern
10. HEBER, Gerät zum Heben
11. RITTER, ein Mensch, der im Mittelalter zur militärischen Adelsklasse gehörte und gute Taten zu vollbringen hatte
12. ETIKETT, ein Stück Papier oder anderes Material, das an etwas befestigt wird; es dient zur Bezeichnung der

- Sache oder zur Angabe ihres Eigentümers oder Bestimmungsorts
13. MUTTER, der weibliche Elternteil
 14. NEST, ein Bau, in dem Vögel Eier legen und ihre Jungen ausbrüten
 15. VERZIERUNG, etwas, das Schönheit geben soll
 16. PFLANZE, ein lebender Organismus, der kein Tier ist
 17. FEDERBETT, eine Bettdecke
 18. WURZEL, der Teil der Pflanze, der in die Erde hineinwächst
 19. SCHLITTSCHUH, ein Rahmen mit einer Klinge, an einem Schuh befestigt; ermöglicht es einer Person, auf dem Eis zu gleiten
 20. BAUM, eine große Pflanze mit hölzernem Stamm
 21. SCHIEDSRICHTER, jemand, der den Ablauf eines Wettspiels regelt
 22. FIRNIS, Flüssigkeit, die dem Holz ein glattes, glänzendes Aussehen verleiht
 23. RAD, kreisförmiger Rahmen, der sich um seine Mitte dreht
 24. JACHT, ein Boot für Vergnügungsreisen
 25. HILFE, Unterstützung
 26. SCHLACHT, Kampf, Krieg
 27. RATSCHLAG, einen Rat bekommen
 28. ERLEDIGUNG, dienstliche Angelegenheit
 29. BEQUEMLICHKEIT, Komfort
 30. ANGST, ein Zustand, in dem jemand Furcht hat
 31. STUFE, eine Stelle in einer Reihenfolge, eine Güte- oder Wertstufe
 32. WÜRDE, Ruhm, Ehre
 33. ZENTIMETER, ein hundertstel Meter
 34. SCHERZ, ein Ausspruch oder Ereignis, das Lachen hervorruft
 35. ABSCHUSS, ein Treffer, der zur Vernichtung des getroffenen Objekts führt
 36. WERK, Arbeit, Arbeitstätigkeit
 37. MASS, Größe
 38. ANZAHL, Gesamtmenge
 39. REIHENFOLGE, Bestimmung, die angibt, welches Ding auf welches folgt
 40. SPIEL, Vergnügung, Sport
 41. FRAGE, eine Äußerung, mittels derer man fragt
 42. REUE, ein Gefühl des Bedauerns
 43. SKALA, eine Folge von Stufen oder eine Abstufung
 44. SPANNUNG, Gefühl der Erregung
 45. DRANG, Antrieb oder Impuls zum Handeln
 46. VERPFLICHTUNG, feierliches Versprechen
 47. WUNSCH, Verlangen, Sehnsucht nach etwas
 48. ERTRAG, Produktion
 49. LÄNGE, Größe eines Abschnitts
 50. INTERVALL, Abstand zwischen Werten oder Qualitätsstufen
 51. VERHÄLTNIS, Wechselbeziehung zweier Größen
 52. ANFANG, der Punkt, von dem an man mißt

Zu jedem der beiden Versuche erhielt jede Vp ein Päckchen von 52 Karten, auf denen die 52 Wörter standen. In jedem Päckchen waren die Wörter in derselben Reihenfolge geordnet, nämlich in der auf Tab. 1 dargestellten Reihenfolge. Dann wurden die Vpn gebeten, die Karten nach der Bedeutungsähnlichkeit der Wörter zu sortieren.

Instruktion: "Vor Ihnen liegt ein Päckchen mit 52 Karten. Auf den Karten stehen Wörter. Ihre Aufgabe ist es, die Kärtchen nach der Bedeutungsähnlichkeit der auf ihnen stehenden Wörter zu sortieren. Da Wörter oft mehr als eine Bedeutung haben, ist jedem Wort eine kurze Definition beigelegt, die seine Bedeutung präzisieren soll. Beachten Sie daher nicht nur die Wörter, sondern auch die Definitionen. Gehen Sie so vor, daß Sie Wörter mit ähnlicher oder nahe verwandter Bedeutung zu einer Gruppe zusammenlegen. In wieviele Gruppen Sie die Kärtchen einteilen, ist gleichgültig."

Zum Schluß der Versuchssitzung wurden die Kärtchen noch einmal ausgegeben, diesmal mit der Instruktion: "Nun zum Schluß sortieren Sie die Kärtchen nochmals, so, wie Sie es zu Anfang gemacht haben. Es geht nicht darum, daß Sie dieselbe Einteilung bilden wie vorhin, sondern versuchen Sie jetzt, die Forderung, die Wörter nach der Ähnlichkeit ihrer Bedeutung zu sortieren, noch besser zu erfüllen. Legen Sie bedeutungsähnliche Wörter zu einer Gruppe zusammen."

4.2.2 Freie Reproduktion und Wiedererkennung

In diesen Versuchen wurde die Gedächtnisleistung der Vpn bei der Reproduktion von Wörtern untersucht, und zwar mit drei Zehnergruppen von Wörtern, die sich im Grad der gegenseitigen Bedeutungsähnlichkeit unterscheiden. Diese drei Wortgruppen wurden als Liste A, B und C bezeichnet. Liste A enthält die Wörter mit sehr hoher Bedeutungsähnlichkeit, während die Wörter auf Liste C einander relativ fern stehen. Die Wörter auf Liste B bilden in dieser Hinsicht eine mittlere Möglichkeit (vgl. auch Tab. 4). Die den Listen entsprechenden Vorlagen A, B und C bestanden jeweils aus einem weißen Blatt Schreibpapier, auf das die 10 Wörter in einer Spalte getippt waren. Die Reihenfolge der Wörter entspricht der auf Tab. 2 angegebenen. Jede Vp hatte nacheinander zwei von diesen Vorlagen zu

lernen, eine unter der Bedingung des freien Reproduzierens, die andere unter der Bedingung des Wiedererkennens. Jeder der beiden Lernversuche lief nach dem folgenden Muster ab:

Die Vp erhielt eine Vorlage, mit der Anweisung, sich die Wörter einzuprägen. Die Darbietung der Vorlage dauerte 10 Sekunden. Es wurden zwei Reproduktionen desselben Materials verlangt, die eine zwei Minuten nach Abschluß der Lernphase, die zweite sieben Minuten nach Abschluß der Lernphase. Im Intervall zwischen dem Lernen und der ersten Reproduktion wurden die Vpn mit der Wiedergabe von vier- bis siebenstelligen Ziffernreihen beschäftigt (analog der Prüfung der Gedächtnisspanne im Wechsler-Test). Den Vpn wurden nacheinander eine Vierer-, eine Fünfer-, eine Sechser- und eine Siebenerreihe von Ziffern diktiert, mit der Anweisung, jede Serie sofort nach dem Diktieren auswendig niederzuschreiben. Danach bekamen die Vpn den Auftrag, weitere Zahlenserien in umgekehrter Reihenfolge niederzuschreiben. Es wurden ihnen wieder zunehmend längere Serien vorgesagt, diesmal mit anderen Ziffern. Im Intervall zwischen der ersten und der zweiten Reproduktion brauchten die Vpn nichts zu tun. Sie wurden gebeten, ruhig auf ihren Plätzen sitzenzubleiben. Sie wußten, daß eine zweite Reproduktion verlangt wurde.

Das Experiment war so angelegt, daß die Hälfte der Vpn (d.h. 21) zuerst unter der Bedingung "freies Reproduzieren" und dann unter der Bedingung "Wiedererkennen" arbeitete, während die andere Hälfte der Vpn die Teilexperimente in der umgekehrten Reihenfolge durchführte. In beiden Teilexperimenten wurde jede der drei Vorlagen A, B und C von je 14 Vpn gelernt (siehe Tab. 2)

Tab. 2: Vorlagen, die in den Versuchen mit freier Reproduktion und mit Wiedererkennen benutzt wurden

Vorlage A	Vorlage B	Vorlage C
Stufe	Anker	Drang
Zentimeter	Koch	Bleichmittel
Maß	Bügeleisen	Arzt
Reihenfolge	Heber	Ritter

Fortsetzung Tab. 2

Skala	Mutter	Schlittschuh
Anzahl	Federbett	Ratschlag
Länge	Schlittschuh	Bequemlichkeit
Intervall	Rad	Angst
Anfang	Jacht	Verpflichtung
Verhältnis	Bequemlichkeit	Stufe

Da der Versuchsverlauf bei beiden Reihenfolgen im wesentlichen gleich war, beschreiben wir nur den Fall, wo das freie Reproduzieren dem Wiedererkennen voranging. Nach der ersten semantischen Distanzmessung bekam jede Vp eine Vorlage mit der Rückseite nach oben. Die Vpn wurden gebeten, den auf der Rückseite ihrer Vorlage stehenden Buchstaben in ein vorbereitetes Formular einzutragen und die Vorlage einstweilen nicht umzudrehen. Dann wurden sie über den Verlauf des Experiments informiert.

Instruktion: "Auf der Vorderseite des Blattes, das Sie vor sich haben, sind einige der Wörter, die Sie vor einer Weile sortiert haben. Jetzt haben Sie die Aufgabe, diese Wörter auswendigzulernen. Die Zeit zum Einprägen ist zwar etwas knapp, aber sie wird gerade ausreichen. Sobald ich es sage, drehen Sie die Vorlage um und fangen an zu lernen. Wenn ich dann "genug" sage, wenden Sie die Vorlage wieder und legen sie mit der Vorderseite nach unten auf den Tisch. Darauf werden wir uns eine Weile mit etwas anderem beschäftigen, und wenn ich es sage, sollen Sie anfangen, die eingepprägten Wörter auf dem Formular zu notieren. Nach einer Pause sollen Sie dann noch einmal an die Wörter denken. Nun, Achtung bitte, drehen Sie die Vorlagen um. Sie können mit dem Lernen beginnen."

Nach Beendigung der Reproduktion wurden die Vorlagen erneut verteilt und zwar nach einem Zufallsplan so, daß jede Vp eine andere Vorlage als zuvor bekam. Es folgte die Instruktion über das Wiedererkennen.

Instruktion: "Dieser Versuch geht wie der vorige mit dem einzigen Unterschied, daß Sie nicht aufschreiben müssen, was Sie sich gemerkt haben, sondern daß Sie zu diesem Versuch das Päckchen Karten, das Sie vorhin sortiert haben, bekommen, um darin die Wörter wiederzufinden, die Sie jetzt lernen sollen."

Bei diesem Versuch wurden die Ergebnisse beider Reproduktionen vom Untersucher selbst notiert.

4.2.3 Paarassoziations-Lernen

Das Lernen von Paarassoziationen folgte auf die beiden soeben beschriebenen Versuche. In diesem Telexperiment hatten die Vpn 18 Wortpaare zu lernen, die aus den 52 Substantiven ausgewählt worden waren. Die Auswahl war so getroffen worden, daß sie es ermöglichte, die Beziehung zwischen der semantischen Distanz der gepaarten Wörter und der Gedächtnisleistung der Vpn zu untersuchen. Tabelle 3 zeigt die Liste der Wortpaare in der Reihenfolge, wie sie im Versuch verwendet wurde.

Tab. 3: Vorlage zum Paarassoziations-Lernen

1. Spannung	- Drang	2. Heber	- Verpflichtung
3. Wurzel	- Baum	4. Zentimeter	- Maß
5. Ritter	- Abschluß	6. Bequemlichkeit	- Reihenfolge
7. Mutter	- Verzierung	8. Schlittschuh	- Schlacht
9. Firnis	- Reue	10. Arzt	- Federbett
11. Fisch	- Nest	12. Erledigung	- Würde
13. Bügeleisen	- Ratschlag	14. Bleichmittel	- Leim
15. Etikett	- Hilfe	16. Anker	- Jacht
17. Stufe	- Anzahl	18. Scherz	- Spiel

Am Versuch nahmen insgesamt 49 Personen teil, das Einprägen dauerte 1 Minute. Die Antworten wurden von den Vpn auf einem speziellen Antwortblatt eingetragen (siehe Abb. 4). Die Reproduktion wurde nicht verzögert, sondern begann sofort nach dem Austeilen der Antwortblätter.

Paarassoziations-Versuch	Datum
Name und Vorname	Geburtsjahr
1. Bleichmittel	
2. Stufe	
3. Wurzel	
4. Ritter	
5. Schlittschuh	
6. Fisch	
7. Erledigung	
8. Firnis	

9. Zentimeter	
10. Anker	
11. Bügeleisen	
12. Arzt	
13. Etikett	
14. Bequemlichkeit	
15. Heber	
16. Mutter	
17. Spannung	
18. Scherz	

Abb. 4: Antwortblatt für das Paarassoziations-Lernen

Instruktion: "Auf der Vorlage, die Sie vor sich haben, stehen Wortpaare. Es ist Ihre Aufgabe, diese Paare zu lernen und sich zu merken, welche Wörter zueinander gehören. Es ist etwas ähnliches, wie wenn man in einer Fremdsprache Vokabeln lernt. Für Ihre Antworten werden Sie dann spezielle Blätter bekommen. Achtung, drehen Sie die Vorlagen um. Sie können mit dem Lernen beginnen."

4.3 Auswertung der semantischen Distanzmessungen

Bei der Verarbeitung der Daten aus beiden Messungen der semantischen Distanzen war zunächst einmal festzustellen, ob die für die Vorlagen der Telexperimente selektierten Wörter tatsächlich in solchen metrischen Beziehungen standen, wie aufgrund der Ergebnisse eines Vorversuchs angenommen wurde. Die Daten sind in dieser Hinsicht zufriedenstellend. Sie werden in späteren Abschnitten detailliert dargestellt.

Die Resultate beider Messungen wurden nach der Methode der hierarchischen Clusteranalyse verarbeitet, und zwar sowohl nach der Maximum- als auch nach der Minimum-Methode. Die Resultate der Messungen sind in den Tabellen 2 und 3 angeführt, die Resultate der Clusteranalysen zeigen die Abbildungen 5 und 6.

Tabelle 2:

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26
2	9																									
3	0	0																								
4	0	0	47																							
5	6	39	1	0																						
6	5	0	2	3	1																					
7	11	46	0	0	36	0																				
8	4	4	1	0	4	17	3																			
9	38	11	0	0	7	0	13	5																		
10	38	10	0	0	7	0	12	5	45																	
11	0	0	43	42	2	2	0	1	0	0																
12	23	11	0	0	8	1	15	7	30	29	0															
13	0	0	34	35	1	6	0	2	1	0	35	0														
14	12	5	1	1	6	15	7	17	12	12	1	13	3													
15	12	6	0	0	5	2	8	9	15	15	0	25	1	9												
16	1	1	1	1	2	25	1	34	1	1	1	1	3	11	3											
17	22	10	0	0	8	0	12	5	29	28	0	32	1	14	23	1										
18	2	1	0	0	3	24	1	35	2	2	1	3	3	14	4	46	2									
19	34	9	0	0	6	0	11	5	46	37	0	25	0	9	19	1	24	1								
20	1	1	1	1	2	25	1	34	1	2	1	2	3	13	3	44	2	46	1							
21	0	0	44	43	2	3	0	1	0	0	40	1	32	1	0	1	0	1	2	1						
22	10	45	0	0	43	0	44	2	12	10	0	13	0	4	7	1	10	1	9	1	0					
23	30	8	0	0	4	0	8	4	32	35	0	24	0	10	14	1	24	2	40	1	1	6				
24	38	7	0	0	6	4	9	6	29	29	0	19	0	11	16	2	21	3	32	1	2	7	29			
25	0	0	1	1	0	0	0	0	0	1	1	0	1	0	2	0	1	0	0	1	1	0	0	0		
26	0	0	0	0	1	1	0	0	0	1	3	0	0	0	3	0	1	0	0	1	0	0	0	0	16	

Fortsetzung Tabelle 2

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

Fortsetzung Tabelle 2

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

Tabelle 2

Resultate der semantischen Distanzmessung vor den Gedächtnisexperimenten (α). Tabelliert sind die Werte N_{ij} . Die Nummern der Zeilen und Spalten entsprechen den Wortnummern in Tabelle 1. Es werden nur die Werte unterhalb der Hauptdiagonalen angegeben.

Tabelle 3:

1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26
2	11																								
3	1	1																							
4	0	1	44																						
5	7	32	1	1																					
6	4	0	8	5	2																				
7	12	44	2	1	30	1																			
8	7	4	2	2	7	16	4																		
9	31	13	1	0	9	0	15	6																	
10	32	13	1	0	10	1	16	6	41																
11	2	0	35	36	1	4	0	2	3	3															
12	24	13	0	0	11	1	16	11	23	24	2														
13	0	0	31	31	1	7	0	2	4	0	29	0													
14	14	7	2	2	8	18	8	23	14	13	3	17	4												
15	16	11	3	0	8	4	9	13	18	16	2	24	6	14											
16	0	0	3	2	2	25	0	32	0	0	1	1	4	16	4										
17	22	11	3	3	10	1	12	9	26	22	2	26	5	18	23	2									
18	0	0	3	2	3	24	0	34	0	0	1	1	4	17	6	49	1								
19	29	9	0	1	10	1	10	7	18	30	1	23	0	14	20	1	21	1							
20	0	0	3	2	3	23	0	34	0	0	1	1	5	18	5	45	1	48	1						
21	0	0	38	38	2	5	0	3	0	0	35	0	28	2	0	2	1	2	9	2					
22	11	46	2	0	32	2	43	4	12	13	0	13	0	6	11	1	11	0	0	0	1				
23	32	9	1	0	12	2	11	9	34	35	2	23	0	14	18	2	21	2	34	2	2	9			
24	33	8	1	0	8	4	9	6	23	25	0	16	0	13	19	1	18	2	34	2	3	7	33		
25	0	1	1	3	0	0	1	0	1	1	1	0	1	0	0	0	2	0	1	0	1	0	0	0	
26	0	0	1	3	1	1	0	1	2	2	10	2	2	2	0	1	2	1	2	1	3	0	1	1	15

Fortsetzung Tabelle 3

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

Fortsetzung Tabelle 3

	27	28	29	30	31	32	33	34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	
28		22																								
29		15	12																							
30		10	9	12																						
31		1	1	1	2																					
32		12	15	21	21	3																				
33		1	1	1	1	41	2																			
34		24	13	14	9	3	11	1																		
35		10	11	6	3	2	6	2	12																	
36		14	27	10	3	6	9	6	8	9																
37		1	2	3	1	38	3	47	2	3	6															
38		1	2	2	1	43	2	42	2	3	6	39														
39		1	2	3	1	43	3	42	2	3	6	40	45													
40		11	18	12	6	3	9	1	26	11	18	2	2	2												
41		34	18	14	6	3	10	5	29	9	15	6	5	6	12											
42		11	10	12	46	2	21	1	9	2	3	1	1	1	7	6										
43		1	1	1	1	46	3	43	1	2	5	41	46	47	2	4	1									
44		8	8	11	44	4	20	3	20	4	3	3	3	3	7	7	43	3								
45		12	11	12	41	3	19	2	12	4	6	2	2	2	8	10	39	2	40							
46		25	26	15	10	6	16	2	20	10	22	3	3	3	17	26	10	2	11	17						
47		13	15	12	38	2	16	1	13	4	4	1	1	1	9	13	38	1	36	41	16					
48		8	14	9	3	5	5	5	6	8	25	5	5	5	10	10	3	4	3	4	12	3				
49		1	1	1	1	33	2	36	1	2	3	37	35	34	2	2	1	35	3	2	2	2	2			
50		1	1	1	1	37	2	36	1	2	3	36	39	40	2	3	1	40	3	2	2	1	2	35		
51		1	1	1	1	35	2	39	1	2	3	38	36	39	2	4	1	39	3	2	2	1	2	35	38	
52		1	1	1	1	35	2	38	1	2	3	36	37	38	2	3	1	38	3	2	2	1	2	37	38	38

Tabelle 3

Resultate der semantischen Distanzmessung nach den Gedächtnisexperimenten (β). Tabelliert sind die Werte N_{ij} . Die Nummern der Zeilen und Spalten entsprechen den Wortnummern in Tabelle 1. Es werden nur die Werte unterhalb der Hauptdiagonalen angegeben.

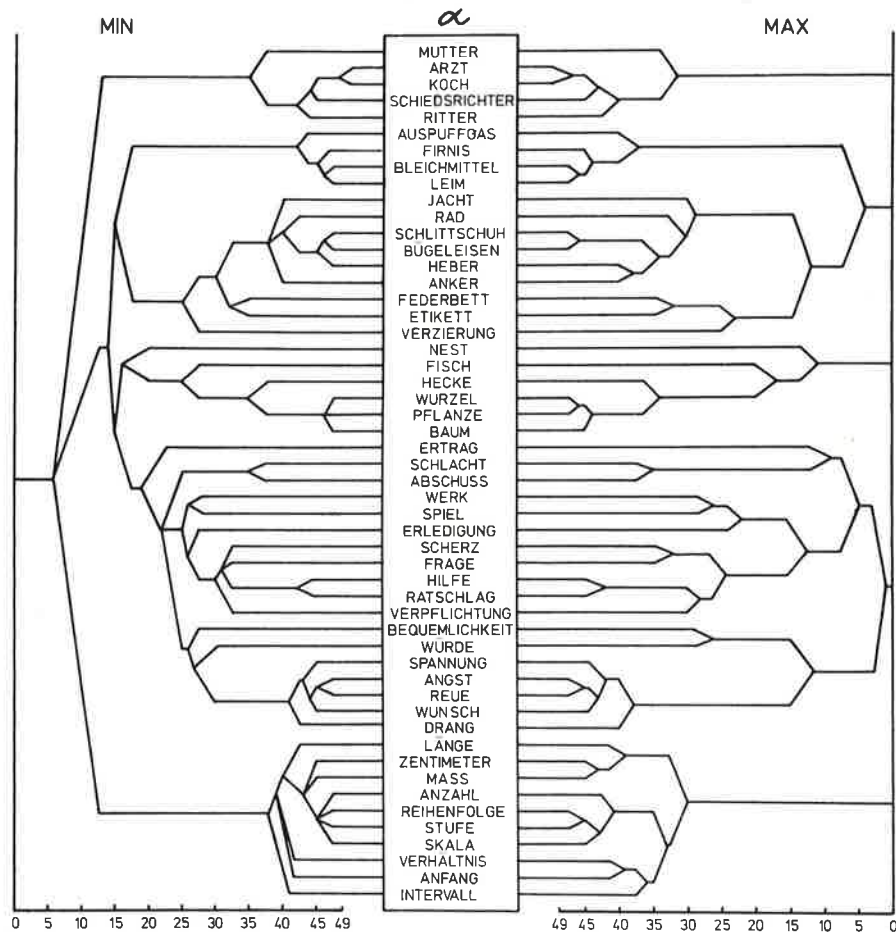


Abb. 5. Ergebnisse der hierarchischen Clusteranalyse,
Daten aus der Messung α

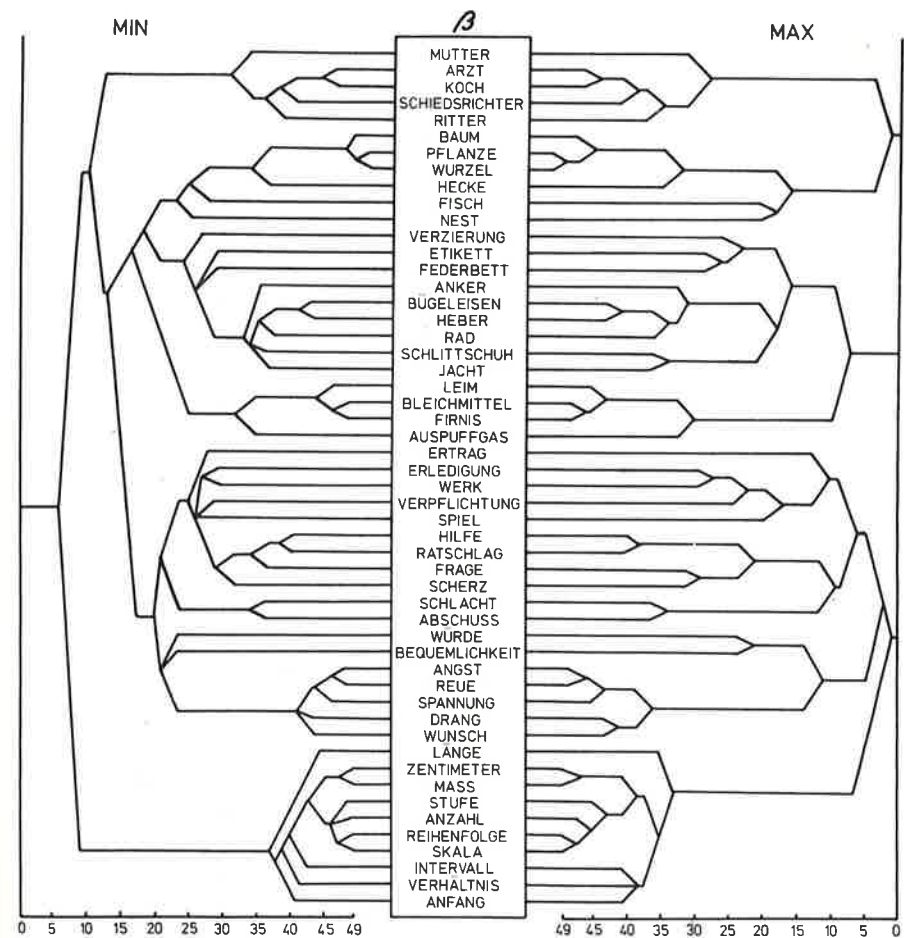


Abb. 6. Ergebnisse der hierarchischen Clusteranalyse,
Daten aus der Messung β

Tabelle 4

Vorlage A

α

	31	33	37	38	39	43	49	50	51	52
31	-	34	38	41	45	43	30	34	34	33
33	34	-	43	39	35	38	39	32	35	35
37	38	43	-	43	41	42	40	35	37	38
38	41	39	43	-	45	42	36	34	37	35
39	45	35	41	45	-	45	33	36	37	35
43	43	38	42	42	45	-	36	38	38	36
49	30	39	40	36	33	36	-	35	39	39
50	34	32	35	34	36	38	35	-	37	36
51	34	35	37	37	37	38	39	37	-	39
52	33	35	38	35	35	36	39	36	39	-

β

	31	33	37	38	39	43	49	50	51	52
31	-	41	38	43	43	42	33	37	35	35
33	41	-	47	42	42	39	36	36	39	38
37	38	47	-	39	40	41	37	36	38	36
38	43	42	39	-	45	46	35	39	36	37
39	43	42	40	45	-	47	34	40	39	38
43	42	39	41	46	47	-	35	40	39	38
49	33	36	37	35	34	35	-	35	35	37
50	37	36	36	39	40	40	35	-	38	38
51	35	39	38	36	39	39	35	38	-	38
52	35	38	36	37	38	38	37	38	38	-

Fortsetzung Tabelle 4

Vorlage B

α

	1	3	9	10	13	17	19	23	24	29
1	-	0	38	38	0	22	34	30	38	1
3	0	-	0	0	34	0	0	0	0	0
9	38	0	-	45	1	29	46	32	29	2
10	38	0	45	-	0	28	37	35	29	2
13	0	34	1	0	-	1	0	0	0	1
17	22	0	29	28	1	-	24	24	21	9
19	34	0	46	37	0	24	-	40	32	1
23	30	0	32	35	0	24	40	-	29	2
24	38	0	29	29	0	21	32	29	-	3
29	1	0	2	2	1	9	1	2	3	-

β

	1	3	9	10	13	17	19	23	24	29
1	-	1	31	32	0	22	29	32	33	2
3	1	-	1	1	31	3	0	1	1	1
9	31	1	-	41	4	26	18	34	23	5
10	32	1	41	-	0	22	30	35	25	2
13	0	31	4	0	-	5	0	0	0	4
17	22	3	26	22	5	-	21	21	18	11
19	29	0	18	30	0	21	-	34	34	2
23	32	1	34	35	0	21	34	-	33	1
24	33	1	23	25	0	18	34	33	-	5
29	2	1	5	2	4	11	2	1	5	-

Fortsetzung Tabelle 4

Vorlage C

 α

	2	4	11	19	27	29	30	31	45	46
2	-	1	0	9	1	1	0	1	0	0
4	1	-	42	0	1	0	0	0	0	0
11	0	42	-	0	1	0	0	0	0	0
19	9	0	0	-	0	1	0	0	0	0
27	1	1	1	0	-	18	14	1	12	28
29	1	0	0	1	18	-	15	0	14	12
30	0	0	0	0	14	15	-	0	40	10
31	1	0	0	0	1	0	0	-	1	3
45	0	0	0	0	12	14	40	1	-	12
46	0	0	0	0	28	12	10	3	12	-

 β

	2	4	11	19	27	29	30	31	45	46
2	-	1	0	9	1	1	0	0	0	0
4	1	-	36	1	1	1	1	0	1	0
11	0	36	-	1	0	1	0	0	0	0
19	9	1	1	-	1	2	1	1	2	1
27	1	1	0	1	-	15	10	1	12	25
29	1	1	1	2	15	-	12	1	12	15
30	0	1	0	1	10	12	-	2	41	10
31	0	0	0	1	1	1	2	-	3	3
45	0	1	0	2	12	12	41	3	-	17
46	0	0	0	1	25	15	10	3	17	-

Tabelle 4

Paarweise semantische Distanzen nach α -Messung und nach β -Messung für die Wörter in den Vorlagen A, B und C. Tabelliert sind die N_{ij} . (Die Numerierung der Zeilen und Spalten in den Submatrizen entspricht den Wortnummern in Tab. 1.)

Weiter wurde für die unter der Hauptdiagonalen liegenden Hälften der beiden Matrizen der Korrelationskoeffizient nach PEARSON berechnet. Für jede der beiden Messungen wurde die prozentuale Übereinstimmung zwischen der Gruppierung nach maximalen Distanzen und der Gruppierung nach minimalen Distanzen berechnet.

Die Korrelation zwischen den beiden Messungen zeigt die Zuverlässigkeit der Messungen an. Sie ist bei Sortierversuchen in der Regel hoch (vgl. z.B. KINTSCH 1970). In unserem Fall ergab sich $r = 0.979$. Das berechtigt uns zu dem Schluß, daß die von den einzelnen Vpn geschätzten semantischen Distanzen durch das summierte Maß hinlänglich genau ausgedrückt werden. Die Mittelwerte der korrelierten Datenmuster betrugen 5.618 und 5.782, die zugehörigen Varianzen 115.542 bzw. 105.110.

Die Ergebnisse der hierarchischen Clusteranalysen werden mit Hilfe von Stammbaumgraphen dargestellt. Für jede Messung wurde eine solche Anordnung der Reizwörter gefunden, daß die der Maximum- und der Minimum-Methode entsprechenden Graphen ohne Linienkreuzungen gezeichnet werden konnten (siehe Abb. 5 und 6). In der graphischen Darstellung der ersten Messung (sie wird im weiteren Text mit α bezeichnet und die zweite Messung mit β), bei der die Maximum-Methode verwendet wurde, finden sich 51 nichtterminale Knoten, im Ergebnis nach der Minimum-Methode sind es 46. Die entsprechenden Anzahlen für die zweite Messung sind 49 bzw. 43. Davon waren bei der ersten Messung den beiden Lösungen 37 Knoten gemeinsam, und bei der zweiten Messung waren es 35 Knoten, d.h. es gab bei beiden Messungen ungefähr 75 % Knoten, die beiden Lösungen gemeinsam waren.

Der Anteil der Knoten, die den beiden Lösungen über denselben Meßdaten gemeinsam sind, zeigt den Grad der Übereinstimmung zwischen den Lösungen an. Wenn die Unterschiede zwischen der Maximum- und der Minimum-Lösung nicht allzu groß sind, können wir, nach JOHNSON (1967), die untersuchte Struktur als hierarchisch ansehen und die Unterschiede zwischen den Lösungen dem Zufall zuschreiben. Genauer wurde diese Frage von MILLER (1969) untersucht. Er hatte analoge Resultate (70 % der Knoten in den Lösungen nach den beiden Methoden stimmten überein) und verglich

sie mit Resultaten einer Monte-Carlo-Simulation. Bei den in der Simulation quasizufällig erzeugten Ähnlichkeitsmatrizen gab es zwischen der Maximum- und der Minimum-Lösung im Durchschnitt nur 14.9 % übereinstimmende Knoten. Dies stützt die Hypothese von der hierarchischen Struktur der Meßdaten.

Es stellt sich aber die Frage nach der Plausibilität der hierarchischen Datenrepräsentation. Plausibel sind die Resultate der Klassifikation dann, wenn sie mit Konzepten der natürlichen Sprache in einer Weise interpretiert werden können, die der Mehrzahl der Sprecher akzeptabel erscheint (vgl. MILLER, S. 181). Schauen wir uns die Abbildungen 5 und 6 an, so finden wir in den Maximum-Lösungen bei den Resultaten der ersten Messung (α) fünf und bei den Resultaten der zweiten Messung (β) drei relativ isolierte Gruppen. Bei den Daten aus der Messung α kann man die Gruppen wie folgt bezeichnen: "Menschen", "Dinge", "Naturalien", "soziale Interaktionsweisen und psychologische Begriffe", "quantitative Begriffe". Die Plausibilität dieser Gruppierung leidet etwas dadurch, daß das Wort "Ertrag" (Produktion) den sozialen Interaktionsweisen und psychologischen Begriffen zugeordnet worden ist. Dieses Wort paßt zu keiner der Gruppen gut, und es ist auch mit der genannten Gruppe nur lose liiert. Bei der Messung β wurden die Vpn durch die Instruktion wohl zu komplizierteren Überlegungen veranlaßt. So erhält man hier die umfassenderen Gruppierungen "Lebewesen und Menschen", "Dinge" und "Abstrakta". In beiden Fällen (also α - und β -Messung) kann man innerhalb der allgemeinsten Gruppen sinnvoll interpretierbare Teilgruppen finden.

Die nach der Minimum-Methode gefundenen Gruppierungen stehen nicht in einer eindeutigen Korrespondenz zu den eben angeführten Befunden. Wie jedoch MILLER in dem zitierten Artikel nachweist, akzentuiert die Minimum-Methode die kleineren und wahrscheinlich unzuverlässigeren Werte N_{ij} , während die Maximum-Methode solche Werte eher unterdrückt. Man kann daher sagen, die Maximum-Methode bildet eine hierarchische Struktur wahrscheinlich zuverlässiger ab. Nichtsdestoweniger bleiben bei beiden Methoden die größeren Distanzen unbestimmt oder werden nur unzuverlässig bestimmt; denn ihre Berechnung basiert nur

auf den Schätzungen von jeweils ganz wenigen Vpn.

4.4 Semantische Distanz und Gedächtnisleistung

4.4.1 Freie Reproduktion und Wiedererkennen

Die Ergebnisse der Versuche mit freier Reproduktion und mit Wiedererkennen sind in Tabelle 5 aufgeführt. Tabelle 5 (F - W) enthält die Ergebnisse der Versuche, in denen das freie Reproduzieren dem Wiedererkennen voranging, Tabelle 5 (W - F) die Ergebnisse der Versuche mit der umgekehrten Reihenfolge. Die Tabellen geben die individuellen Resultate aus beiden Teillexperimenten wieder und zeigen außerdem die Leistungen im RAVEN-Test und in dem hier nicht weiter verarbeiteten verbalen Gedächtnistest. Zu jeder Reproduktion ist die jeweilige Vorlage angegeben. Es werden zwei Angaben über die Leistung bei jeder der beiden Reproduktionen nach 2 bzw. nach 7 Minuten gemacht, nämlich die Anzahl der richtig reproduzierten Wörter und die Anzahl der fälschlich erinnerten Wörter, die sich nicht auf der Vorlage befanden.

In den beiden Tabellen sind auch die mittleren Leistungen unter den einzelnen Versuchsbedingungen aufgeführt. Um einen möglichen Reihenfolgeeffekt der Reproduktionsbedingungen zu überprüfen, wurden korrespondierende Mittelwerte aus Tabelle F - W und Tabelle W - F miteinander verglichen. Die folgenden Vergleiche wurden mit t-Tests überprüft: freie Reproduktion I (6.05 und 6.38), freie Reproduktion II (6.33 und 6.2), Wiedererkennen I (8.10 und 7.08), Wiedererkennen II (8.15 und 7.43). Da alle diese Unterschiede nicht signifikant wurden, konnten die Daten in zwei komprimierteren Tabellen zusammengefaßt werden (siehe Tab. 6). Die erste Teiltabelle zeigt die Ergebnisse des freien Reproduzierens, die zweite die des Wiedererkennens, jeweils für jede der drei Vorlagen (A, B, C).

Tabelle 5 (F - W)

	IQ	P	Freies Reproduzieren				Wiedererkennen			
			I.r.	I.f.	II.r.	II.f.	I.r.	I.f.	II.r.	II.f.
1	104	42	B 5	0	6	1	C 7	1	6	4
2	92	32	A 4	2	4	2	C 5	1	5	2
3	106	25	C 3	3	4	3	B 10	4	7	5
4	112	45	B 7	0	8	1	A 10	0	10	0
5	108	28	A 4	2	3	2	C 6	3	8	2
6	124	29	B 8	0	8	0	A 10	0	10	0
7	108	49	C 5	1	5	1	A 10	0	10	0
8	130	34	A 8	0	9	0	B 9	0	9	0
9	126	39	C 6	0	5	1	B 8	0	9	0
10	102	50	C 4	2	5	2	A 10	0	10	0
11	106	29	B 6	0	7	0	C 7	3	6	5
12	120	52	A 7	0	7	1	B 10	0	10	0
13	106	47	C 9	0	8	0	A 10	0	9	1
14	118	46	A 9	0	9	0	B 9	1	9	1
15	130	44	B 8	0	10	0	A 9	0	10	1
16	122	49	A 5	0	5	1	B 7	0	8	0
17	126	50	B 9	0	9	0	C 9	1	9	1
18	126	40	A 6	0	7	1	C 4	2	4	2
19	116	57	B 8	0	8	0	C 7	3	8	2
20	114	34	C 4	1	4	1	A 8	0	8	0
21	98	29	C 2	1	2	1	B 5	2	6	2
$\bar{x}$	114	40,5	6,05	0,57	6,33	0,857	8,10	1,0	8,15	1,33
s^2	120,7	89,45	4,45	0,86	5,036	0,73	2,69	1,7	1,85	2,635

Fortsetzung Tabelle 5 (W - F)

	IQ	P	Wiedererkennen				Freies Reproduzieren			
			I.r.	I.f.	II.r.	II.f.	I.r.	I.f.	II.r.	II.f.
1	116	49	C 7	1	6	3	B 7	1	6	1
2	112	40	B 9	2	7	2	C 8	1	8	1
3	114	25	A 9	1	9	3	B 5	1	5	1
4	122	51	A 7	0	7	0	C 5	1	5	1
5	116	40	C 4	4	3	7	B 7	2	9	1
6	114	38	B 8	3	7	3	A 7	0	7	0
7	110	34	A 10	0	10	0	C 4	2	2	2
8	124	42	B 8	1	8	1	C 5	2	6	1
9	122	38	C 6	4	4	6	A 7	1	6	1
10	130	53	A 10	0	10	0	C 9	0	9	0
11	124	64	C 7	4	8	0	A 8	1	7	1
12	128	38	B 6	0	6	0	A 8	0	8	1
13	122	38	C 8	2	8	2	B 9	0	9	0
14	120	46	A 8	0	8	0	B 10	0	10	0
15	118	57	A 8	0	9	1	C 7	1	7	1
16	118	38	B 8	0	8	0	A 7	0	6	0
17	120	42	C 4	6	2	8	B 2	4	2	2
18	120	48	B 10	1	9	1	C 2	1	3	1
19	122	49	C 9	1	9	1	A 6	2	4	2
20	122	47	A 10	0	10	0	B 6	0	6	0
21	128	49	B 8	0	8	0	A 5	0	5	0
$\bar{x}$	120	48,8	7,8	1,43	7,43	1,81	6,38	0,95	6,2	0,81
s^2	28,8	74,19	3,06	3,16	9,86	5,96	4,6	1,05	5,06	0,51

Tab. 5. Resultate der Versuche mit freiem Reproduzieren und Wiedererkennen (individuelle Daten)

Erläuterungen:

F-W: das freie Reproduzieren ging dem Wiedererkennen voran

F-W: das Wiedererkennen ging dem freien Reproduzieren voran

Die Vpn wurden für die Teile F-W und W-F getrennt numeriert

IQ: Intelligenzquotient nach RAVEN

P: Rohwerte beim verbalen Gedächtnistest
A, B, C: die Vorlagen
r: Anzahl der richtig reproduzierten Wörter
f: Anzahl falscher Reproduktionen
I: erste Reproduktion
II: zweite Reproduktion

Tabelle 6

Freies Reproduzieren

Vorlage	Mittelwert Varianz	I.r.	I.f.	II.r.	II.f.
A	$\bar{x}$	6,50	0,57	6,20	0,86
	s^2	2,50	-	3,41	-
B	$\bar{x}$	6,90	0,57	7,35	0,50
	s^2	4,23	-	4,87	-
C	$\bar{x}$	5,20	1,14	5,20	1,14
	s^2	4,87	-	4,95	-

Wiedererkennen

Vorlage	Mittelwert Varianz	I.r	I.f.	II.r.	II.f
A	$\bar{x}$	9,20	0,07	9,30	0,43
	s^2	1,10	-	1,00	-
B	$\bar{x}$	8,20	1,00	7,93	1,07
	s^2	2,18	-	9,15	-
C	$\bar{x}$	6,43	2,57	6,15	3,21
	s^2	2,88	-	5,06	-

Tab. 6. Ergebnisse der Versuche mit freiem Reproduzieren und Wiedererkennen (Mittelwerte und Varianzen bei den Vorlagen A, B und C)

Erläuterungen:
r: Anzahl der richtig reproduzierten Wörter
f: Anzahl falscher Reproduktionen
I: erste Reproduktion
II: zweite Reproduktion

Diese zusammengefaßten Werte wurden nun hinsichtlich der folgenden Hypothese analysiert: Vorlage A, die semantisch ähnliche Wörter enthält, wird besser reproduziert als Vorlage B und diese wiederum besser als Vorlage C, deren Wörter ziemlich hohe paarweise semantische Distanzen haben. Die paarweisen semantischen Distanzen der auf den einzelnen Vorlagen verwendeten Wörter sind in Tabelle 4 angegeben. Die Mittelwerte der Distanzen zeigt Tabelle 7.

	A	B	C
α	44,27	16,38	5,23
β	38,71	15,00	5,38

Tab. 7. Mittelwerte der semantischen Distanzen auf Liste A, B, C

Für die Mittelwerte in Tabelle 6 wurden 4 einfaktorielle Varianzanalysen über die Vorlagen gerechnet, und zwar je eine für die freie Reproduktion I und II und das Wiedererkennen I und II (abhängige Variable: richtige Reproduktionen, Spalten r.). Da jede der drei Vorlagen von 14 Vpn bearbeitet worden ist, bezieht sich jede Varianzanalyse auf eine 3x14-Tabelle. (Weil die Analysen auf "repeated measurements" basieren, sind die F-Tests von Reproduktion I und II nicht unabhängig.) Es ergaben sich die folgenden F-Werte:
Freie Reproduktion I: F = 2,83, II: F = 3,65
Wiedererkennen I: F = 13,53, II: F = 6,87.

Unter der Nullhypothese "die Leistungen bei den unterschiedlichen Vorlagen unterscheiden sich nicht" haben die gefundenen F-Werte mit 2 und 39 Freiheitsgraden die folgenden Wahrscheinlichkeiten p (jeweils Wahrscheinlichkeit über dem Intervall oberhalb des gefundenen Werts):

für $F = 2,86$ ist p größer als 10 %,

für $F = 3,65$ ist p kleiner als 5 %,

für $F = 6,87$ ist p kleiner als 1%.

Wir sehen, daß die Hypothese für das freie Reproduzieren II auf dem 5%-Signifikanzniveau und für das Wiedererkennen I und II auf dem 1%-Signifikanzniveau gestützt ist. Unter allen vier Bedingungen zeigte sich die erwartete Reihe der Leistungen bei den drei Vorlagen übereinstimmend: Mittelwert bei A höher als bei B und bei B höher als bei C.

4.4.2 Paarassoziations-Lernen

Zur Analyse der Ergebnisse des Paarassoziations-Lernens haben wir die Reproduktionsleistung bei jedem assoziierten Wortpaar durch die Anzahl der Vpn (von 49) ausgedrückt, die das jeweilige Reizwort mit dem richtigen Reaktionswort beantwortet haben. Die Häufigkeiten richtiger Antworten auf die Reaktionswörter haben wir mit den semantischen Distanzen der zugehörigen gepaarten Wörter verglichen. Diese Daten sind in Tabelle 8 angegeben.

Tab. 8. Ergebnisse des Paarassoziations-Versuchs

α : Distanz nach α -Messung, } tabelliert sind N_{ij}
 β : Distanz nach β -Messung }
 r: Leistung, f: Fehlerzahl

Wortpaare	α	β	r	f
1. Spannung - Drang	41	40	29	6
2. Heber - Verpflichtung	1	1	16	11
3. Wurzel - Baum	46	48	41	0
4. Zentimeter - Maß	43	47	42	1

5. Ritter - Abschluß	2	9	39	7
6. Bequeml. - Reihenfolge	1	3	14	6
7. Mutter - Verzierung	1	6	35	3
8. Schlittschuhe - Schlacht	0	2	21	3
9. Firnis - Reue	0	0	11	3
10. Arzt - Federbett	0	3	32	4
11. Fisch - Nest	15	18	32	1
12. Erledigung - Würde	21	15	24	4
13. Bügeleisen - Ratschlag	0	0	16	5
14. Bleichmittel - Leim	46	44	36	9
15. Etikett - Hilfe	0	0	15	6
16. Anker - Jacht	38	33	39	3
17. Stufe - Anzahl	41	43	10	11
18. Scherz - Spiel	19	26	41	0

Die Beziehung zwischen Reproduktionshäufigkeiten und Distanzen wird durch den Rangkorrelationskoeffizienten nach SPEARMAN ausgedrückt. Der Koeffizient beträgt für die α -Messung $\rho = -0.514$ und für die β -Messung $\rho = -0.623$. Es ist offensichtlich, je größer die semantische Distanz der assoziierten Wörter, desto kleiner ist die Wahrscheinlichkeit der Reproduktion der richtigen Assoziation. Die Korrelation mit der β -Messung ist vermutlich deswegen höher, weil durch das Assoziieren sich gewisse Paare angenähert haben, was sich dann in der β -Messung ausdrückt. D.h. zu der in der α -Messung festgestellten Nähe kommt jetzt noch eine Annäherung durch die erfolgreiche Assoziation (und Reproduktion) hinzu (vgl. Tab. 4).

Die Tabelle 5, 6 und 8 enthalten außer den Daten über die Leistungen der Vpn auch Angaben über die Anzahl der Fehler, d.h. der Worтеinfälle, die nicht auf den Vorlagen dargeboten worden waren. In allen Versuchen richtete sich die Fehlerzahl offensichtlich nach ähnlichen Gesetzmäßigkeiten wie die Behaltensleistung, nur daß die Beziehung der Fehlerzahl zu den semantischen Distanzen in Wortpaaren schwächer ist und das umgekehrte Vorzeichen hat. Am besten sieht man das an den Ergebnissen des Paarassoziations-Lernens. (Rangkorrelationskoeffizient nach SPEARMAN zwischen Fehlerzahl und semantischer Distanz nach α -Messung $\rho = 0.073$ und nach β -Messung $\rho = 0.206$).

4.5 Diskussion

In den beschriebenen Experimenten wurden drei miteinander zusammenhängende Hypothesen überprüft. Die erste betrifft die Beziehung des semantischen Distanzmaßes zu den gemeinsamen semantischen Merkmalen der Wortpaare, die zweite betrifft gewisse Eigenschaften dieses Maßes und die dritte die Beziehung zwischen der Gedächtnisleistung und der semantischen Distanz der eingeprägten Wörter.

1. Die erste Hypothese besteht aus zwei Annahmen: Nach MILLER (1969) gelangen die Wörter aufgrund gemeinsamer Merkmale in eine gemeinsame Gruppe, und zwar werden zwei Wörter desto häufiger in eine gemeinsame Gruppe sortiert, je mehr gemeinsame Merkmale sie haben. Wir wiederum haben angenommen, daß man sich ein Wortpaar desto leichter merken kann, je mehr Merkmale es gemeinsam hat. Die Konjunktion beider Annahmen impliziert, daß die Behaltensleistung für Wortpaare desto höher sein sollte, je geringer ihre nach MILLER definierte Distanz ist. Die Resultate vom Paarassoziations-Lernen entsprechen dieser Erwartung. Ergo ist unsere Konzeption, daß die semantische Distanz zweier Namen in umgekehrter Beziehung zur Anzahl der ihnen gemeinsamen semantischen Merkmale steht, plausibel.

2. Die semantische Distanzmessung hatte zwar im Experiment nur eine Hilfsfunktion; es ist aber für die weitere Argumentation nicht unwichtig zu wissen, ob es sich bei dem Maß um eine Ultrametrik handelt. In der Tat ist der Grad der Übereinstimmung der beiden hierarchischen Gruppierungsverfahren so hoch und die gefundenen Gruppen lassen sich so gut interpretieren, daß man eine Ultrametrik annehmen kann.

3. Die Versuche, in denen wir die Beziehung zwischen der Behaltensleistung und der semantischen Distanz der behaltenen Wörter untersucht haben, sollten letztlich zur Entscheidung beitragen, ob die Struktur des semantischen Gedächtnisses hierarchisch oder paradigmatisch ist. Sofern die Struktur hierarchisch ist, gilt, wie wir früher erwähnt haben, eine der in Hypothese 1 über Wortpaare gemachten Annahme entsprechende Beziehung auch für k -Tupel von Wörtern, mit $k > 2$: Von

einer Wortmenge wird desto mehr reproduziert, je "kompakter" sie ist, d.h. geringer die Größenordnung der paarweisen semantischen Distanzen in ihr. Die Resultate des Wiedererkennungsversuchs stützen nun diese Hypothese. Die Leistung beim Wiedererkennen fällt nämlich monoton von Vorlage A zu Vorlage C ab. Weniger deutlich zeigen sich diese Verhältnisse an den Daten vom freien Reproduzieren. Zwar ist die Leistung bei Vorlage C auch am schwächsten, aber der Abfall ist nicht ganz monoton. Zusammengenommen ergeben die Resultate der Gedächtnisexperimente ein ziemlich starkes Argument zugunsten der Annahme einer hierarchischen Struktur.

5. DER ZUSAMMENHANG EINER HIERARCHISCHEN GEDÄCHTNISSTRUKTUR MIT EINEM INDUKTIV ARBEITENDEN ERKENNUNGSMODELL

Eine hierarchische Gedächtnisstruktur, deren Annahme durch die im vorigen Kapitel berichteten Experimente gestützt wird, entspricht einem hierarchischen Erkennungsmodell, wie wir es im 2. Kapitel (S. 118) besprochen haben. Die Experimente scheinen also auch ein hierarchisches Modell des Erkennens zu fordern. Falls man darunter jedoch ein seriell arbeitendes Modell versteht, das die Merkmale des zu erkennenden Objekts nacheinander anhand eines gespeicherten hierarchischen Katalogs entscheidet, so gerät man damit in gewisse Schwierigkeiten. Faktisch lassen sich nämlich die identifizierten Objekte nicht immer exakt in eine Hierarchie einordnen (WATANABE 1965). Beim Einprägen neuer Items wäre öfters eine Umklassifikation sämtlicher schon gespeicherten Items erforderlich. Aus diesen Gründen arbeitet man ja auch im Bibliothekswesen nicht ausschließlich mit hierarchischen Katalogen, sondern ergänzt sie durch Systeme, die eher unserem parallelen Modell entsprechen (MEADOW 1967). Die aus dem Bibliothekswesen bekannten Schwierigkeiten mit hierarchischen Katalogen scheint es nun aber beim Gedächtnis nicht zu geben. Also hat man sich zu fragen, warum dies nicht der Fall ist. Entweder sind die experimentellen Belege für eine hierarchische Gedächtnisstruktur

tur doch nicht so schlüssig, oder die Analogie mit einer Bibliothek trifft nicht ganz, d.h. das Gedächtnis hat zwar eine hierarchische Struktur, aber sie funktioniert anders, als man es von einem hierarchischen Bibliothekskatalog kennt. Gäbe es als Beleg für eine hierarchische Struktur des Gedächtnisses nur die im vorigen Kapitel beschriebenen Experimente, so wäre das kein ausreichender Beweis. In der Literatur finden wir aber eine Reihe von Belegen, die denselben Schluß wie unser Experiment nahelegen.

So wurde z.B. in Versuchen zur Reproduktion von gruppenweise dargebotenen Wörtern (HALMIOVÁ & ŠIPOŠ 1976) wiederholt festgestellt, daß Wörter, die zu einer und derselben Bedeutungskategorie gehören, leichter und besser reproduziert werden als Wörter von unterschiedlicher Bedeutung. KINTSCH (1970) führte ein Experiment durch, das den in dieser Arbeit beschriebenen Versuchen in methodischer Hinsicht nahesteht. Er maß die semantischen Distanzen von Wörtern mit derselben Methode wie MILLER (1967). Die verwendete Wörtermenge ließ er von Vpn auswendig lernen, und dann beobachtete er die Reihenfolge der Wörter beim freien Reproduzieren. Er fand, daß sukzessiv reproduzierte Wörter "chunks" von relativ hoher Bedeutungsnähe bildeten. In einem weiteren Versuch ließ KINTSCH die Vpn ein aus derselben, ihnen bekannten Wortmenge ausgewähltes Wort mit Hilfe von 20 Ja-Nein-Fragen erraten. Die Aufgabe war eine Variante des bekannten Gesellschaftsspiels "Zwanzig Fragen". Wie KINTSCH zeigen konnte, bewegten sich die Vpn dabei im wesentlichen längs den Ästen eines hierarchischen Graphen, der aus der semantischen Ähnlichkeitsmatrix der Wörter mit JOHNSON's hierarchischer Clustermethode konstruiert worden war (1967).

Die angeführten Belege für die Hierarchizität des Gedächtnisses stammen alle aus Laborexperimenten. Unter natürlichen Bedingungen versuchte ich die Annahme in dem folgenden kleinen Experiment zu überprüfen. Nach einem längerdauernden Kursus für Bauleiter (20 - 45 Jahre alt, meist Fachschulabschluß) bat ich die Teilnehmer, sich an den Inhalt aller meiner Vorlesungen zu erinnern und die einzelnen Items in der Reihenfolge, wie sie

ihnen einfielen, zu notieren. Das Experiment wurde in drei unabhängigen kleinen Stichproben (16, 14 bzw. 17 Vpn) durchgeführt (KOLMAN 1977). Dabei hatte ich die folgende Hypothese im Sinn:

Wenn das Gedächtnismaterial semantisch klassifiziert ist, so wird es besser behalten und leichter reproduziert. Daher sollten diejenigen Vpn relativ viele Items reproduzieren können, die relativ häufig bedeutungsnahe Items unmittelbar nacheinander reproduzieren.

Für jede Vp wurde die Gesamtzahl der Wörter, die Gesamtzahl der erinnerten Items und die Anzahl der mit dem unmittelbaren Vorgänger in der Reproduktionsreihe inhaltlich zusammenhängenden Items ermittelt. Den letztgenannten Wert nennen wir im folgenden der Einfachheit halber "Anzahl der Beziehungen". Als eine Beziehung wurde gewertet, wenn das gegebene Item zu demselben thematischen Kreis wie das vorhergehende Item gehörte (z.B. beide betreffen Organisation und Management bzw. Sozialpsychologie bzw. Arbeitsschutz etc.). Die semantischen Beziehungen waren also in einer inhaltlichen hierarchischen Klassifikation verankert.

Für jede Stichprobe wurden zwei Rangkorrelationskoeffizienten nach SPEARMAN berechnet, und zwar der eine zwischen der Zahl der Beziehungen und der Zahl der Wörter und der andere zwischen der Zahl der Beziehungen und der Zahl der Items. Tabelle 9 zeigt diese Koeffizienten.

Tab. 9. Korrelationen zwischen der Menge reproduzierten Materials und der Anzahl semantischer Bindungen bei sukzessiv reproduzierten Items

Stichprobe:	A	B	C
Bindungen - Wörter	0.52	0.39	0.92
Bindungen - Items	0.83	0.85	0.93

Trotz der kleinen Stichproben sind alle Koeffizienten außer dem von B in der ersten Zeile statistisch signifikant. Die Zuverlässigkeit des Resultats zeigt sich darin, daß es in drei unabhängigen Stichproben übereinstimmend gefunden

wurde.

Die Resultate des Experiments stützen die getestete Hypothese. Die Korrelationskoeffizienten belegen, daß die Vpn desto mehr Items reproduziert haben und daß ihre Reproduktionen desto mehr Wörter enthielten, je mehr die Vpn sich bei der Wiedergabe an die semantischen Beziehungen im Material hielten. Daß die Korrelationen in der zweiten Zeile von Tabelle 8 durchweg höher als die in der ersten Zeile sind, zeigt, daß die Ausnutzung der semantischen Bindungen im eingepprägten Material nicht bloß die Wortproduktion steigert, sondern in erster Linie den wiedergegebenen Inhalt vermehrt. Dies Ergebnis stimmt ebenfalls gut mit unserer Auffassung von der Wirkungsweise eines hierarchisch geordneten Speichers überein.

5.1 Diskriminationsstärke von Erkennungsmodellen

Wie es scheint, sind die Argumente für ein hierarchisches Gedächtnis hinreichend stark. Demnach ist das Gedächtnis entsprechend einem hierarchischen Katalog organisiert. Nun ist aber, wie wir schon angedeutet haben, die Benutzung eines hierarchischen Katalogs mit gewissen Schwierigkeiten verbunden. Jedoch bei menschlichen Gedächtnisleistungen treten diese Schwierigkeiten nur selten auf, ja sie sind ganz und gar untypisch. Wir wollen versuchen zu zeigen, worin sich ein hierarchisches menschliches Gedächtnis von einem hierarchischen Bibliothekskatalog unterscheiden muß. Zu dem Zweck führen wir einen neuen Begriff ein: Diskriminationsstärke. Wir wollen die Diskriminationsstärken des hierarchischen und des parallelen Erkennungsmodells vergleichen.

Was Diskriminationsstärke heißt, zeigen wir am Beispiel des Gesellschaftsspiels, das KINTSCH (1970) in dem oben angeführten Experiment verwendet hat. Dieses Spiel heißt in der angelsächsischen Literatur "Twenty Questions". Es geht so: Ein Spieler soll mit höchstens zwanzig Fragen einen Gegenstand erraten, den die anderen Spieler vorher vereinbart haben. Auf seine Fragen wird nur mit Ja oder Nein geantwortet. Wenn die Fragen gut gewählt sind, so daß durch jede die Menge der mög-

lichen Dinge ungefähr um die Hälfte reduziert wird, so kann man mit 20 Fragen ein Ding aus 2^{20} möglichen Dingen herausfinden. Dies ist die Diskriminationsstärke eines Systems von 20 gut gewählten Fragen.

Die Identifizierung eines Gegenstandes in diesem Gesellschaftsspiel durch 20 nacheinander gestellte Fragen erinnert an die Zuordnung eines Konzepts (Namen) zu einem Muster gemäß dem hierarchischen Modell. Der Fragende könnte natürlich auch wie ein Parallelmodell verfahren. In dem Fall müßte er erst alle Fragen aussprechen und den Eingang aller Antworten abwarten, ehe er mit der Bestimmung des gesuchten Dinges beginnt. Nehmen wir an, der Fragende hätte im voraus zwanzig geeignete Fragen gewählt und sorgfältig vorbereitet. Anstatt nun eine nach der anderen zu stellen und nach jeder Antwort seinen Suchbereich einzuengen, gibt er den Mitspielern alle Fragen auf einmal, gleichsam wie einen Fragebogen. Er tut also dasselbe wie ein Parallelmodell, das auf ein dargebotenes Muster gleichzeitig 20 Charakterisatoren anwendet. Kann dies parallele Fragenstellen genauso effektiv sein, d.h. genauso oft zum Erfolg führen wie das übliche sequentielle Fragenstellen? Wer das Spiel kennt, wird sicherlich Nein antworten. Dies ist zwar nur ein intuitiver Schluß, aber das bedeutet nicht, daß er keinen rationalen Kern hat.

Wir wollen nun durch eine kombinatorische Überlegung die Diskriminationsstärken von parallelem und hierarchischem Erkennungsmodell vergleichen. Legen wir eine bestimmte Anzahl von binären Dimensionen, sagen wir n , fest. Mit ihrer Hilfe können 2^n Muster beschrieben werden, die sich paarweise in mindestens einem Merkmal unterscheiden. Jeder Dimension soll ein Charakterisator entsprechen. Wir ordnen die Charakterisatoren zu einem Parallelmodell an. Dieses Modell kann dann mit n Charakterisatoren genau 2^n Muster identifizieren. In einem hierarchischen Modell wäre die Situation komplizierter. Um 2^n Muster identifizieren zu können, muß man aus Charakterisatoren einen Stammbaum mit 2^n Endknoten bilden. Der komplette binäre Entscheidungsbaum enthält 2^n Äste (genauer: Wege). Auf jedem Ast liegen genau n Knoten, die je einem Cha-

rakterisator entsprechen. Zwar haben die verschiedenen Wege durch den Graphen paarweise unterschiedlich lange Abschnitte gemeinsam, doch muß die Gesamtzahl der für die Diskrimination von 2^n Mustern erforderlichen Charakterisatoren im hierarchischen Modell größer sein als im parallelen. Diese Anzahl beträgt

$$\sum_{i=0}^{n-1} 2^i \text{ und ist somit größer als } n.$$

Der Vergleich scheint also deutlich zugunsten des Parallelmodells auszugehen. Es ist einfacher strukturiert, und es leistet dieselbe Diskrimination mit weniger Charakterisatoren. Dabei muß seine Identifizierungsgeschwindigkeit nicht einmal geringer sein als die des hierarchischen Modells. Was nun aber die semantische Struktur des menschlichen Gedächtnisses betrifft, so zeigen die bisherigen Befunde relativ eindeutig, daß sie hierarchisch ist. Das bedeutet, daß der Mensch zum Erkennen ein hierarchisches Modell befolgt. Dennoch ist das menschliche Erkennen sehr effektiv, entschieden effektiver als alle derzeit auf Computern programmierbaren Identifikationsleistungen, und das trotz der enormen Arbeitsgeschwindigkeit der Computer. Und es ist eigentlich unplausibel, daß die Prinzipien unserer Diskriminationsfähigkeit ohne zwingenden Grund so offensichtlich ungünstig sind, wie es eben für das hierarchische Modell gezeigt wurde.

Vergleichen wir einmal die kombinatorische Überlegung, die offenbar für das Parallelmodell spricht, genauer mit unseren Erfahrungen beim Zwanzig-Fragen-Spiel, so bemerken wir, daß sie nicht ganz in Ordnung ist. Die kombinatorische Betrachtung der beiden Modelle setzt nämlich voraus, daß alle Merkmalsdimensionen, also auch die Charakterisatoren, fest vorgegeben sind. Ein Fragender aber, der das Zwanzig-Fragen-Spiel ernsthaft spielen will, weiß nicht im voraus, welche Fragen er stellen wird, ja, er wäre nicht einmal in der Lage, sie im voraus festzulegen. Unsere kombinatorische Überlegung ging aus von einer Auffassung des Erkennens, die dieses sozusagen als einen deduktiven Prozeß sieht. Das Fragenstellen im Zwanzig-Fragen-Spiel ist aber ein induktiver Prozeß.

Nach einer deduktiven Auffassung des Erkennens ist die

Zahl der identifizierbaren, d.h. benennbaren Muster kombinatorisch durch die Zahl der verfügbaren Merkmalsdimensionen bestimmt. Eine Hierarchisierung der Dimensionen erleichtert das Auffinden eines Musternamens nicht, sondern verkompliziert es eher, da die Hierarchie redundante Dimensionen enthält. Wenn wir nun die oben erwähnte Diskriminationsstärke eines Erkennungsmodells definieren als den Quotienten aus der Anzahl Charakterisatoren, mit denen das Modell arbeitet, und dem Umfang der Namenmenge, die den erkennbaren Mustern entspricht, so hat beim deduktiven Erkennen das Parallelmodell die größere Diskriminationsstärke. Wenn jedoch das Erkennen ein induktiver Vorgang ist, dann ist die Hierarchisierung der Merkmalssuche die beste Art, wie das Aufsuchen der Namen optimiert werden kann. Ganz allgemein ist natürlich die Diskriminationsstärke eines induktiven Modells größer als die eines deduktiven Modells, das gleich viele Unterscheidungen treffen kann. Wir können das wieder am Zwanzig-Fragen-Spiel demonstrieren. Wir haben gesagt, daß 20 gut gewählte Fragen ein Element in einer Menge von 2^{20} möglichen Elementen identifizieren können. Aber der fragende Spieler weiß nicht im voraus, in was für einer Menge der Zielgegenstand zu suchen ist. Für den Spieler steht die Menge von 2^{20} Möglichkeiten am Ende und nicht am Anfang seines Fragens. In der Tat hat er den Zielgegenstand in einem anfangs unbestimmten Universum von Gegenständen aufzufinden. Um dabei Erfolg zu haben, muß er seine Fragen sukzessiv erzeugen und die Antworten auf die vorausgehenden Fragen berücksichtigen.

Die gegenwärtige theoretische Auffassung des Mustererkennens versteht das Erkennen als eine Deduktion aus den hier als Prämissen fungierenden, schon am Muster konstatierten Merkmalen. Betrachten wir das Erkennen jedoch als eine Induktion analog dem Fragenstellen in dem eben beschriebenen Gesellschaftsspiel, dann erscheint eine Reihe der hier diskutierten Probleme in einem neuen Licht. Eine Induktion kann überhaupt nur als ein sukzessives Stellen von Fragen ablaufen, und hierarchische Suche und hierarchische Struktur werden zur Notwendigkeit. Auch ist bei einem induktiven hierarchischen Modell

die Fehlerlosigkeit der Identifikation nicht so sehr durch etwaige Unvollkommenheit der Charakterisatoren bedroht wie in dem vorhin diskutierten deduktiven Modell. Eine unzweckmäßig gewählte Frage muß noch nicht bedeuten, daß die Induktion ihr Ziel verfehlt. Andererseits gibt es bei einem induktiven Prozeß keine Garantie, daß er immer zur Wahl des richtigen Namens, ja überhaupt zur Wahl eines Namens führt.

5.2 Induktives Erkennen

Die Vorstellung, daß das Identifizieren ein induktiver Prozeß sei, ist nicht neu. Einige Autoren modellieren den Erkennungsprozeß mit Hilfe von Problemlöseprogrammen (UHR 1973: 189). ZINČENKO & VERGILES (1972) betonen die produktiven, schöpferischen Aspekte der Wahrnehmung. Die induktive hierarchische Merkmalssuche ermöglicht die Identifikation von praktisch unbegrenzt vielen unterschiedlichen Objekten. Im Laufe seines Lebens identifiziert ein Mensch sicherlich eine sehr große Menge von Dingen, sei es als individuelle Dinge (z.B. den eigenen Mantel unter den Mänteln in einer Garderobe), sei es als Elemente von Klassen (z.B. ein Auto einer bestimmten Marke). Obwohl ein Mensch während seines Lebens nur eine endliche Menge von Dingen identifiziert, bildet doch die Gesamtheit der Dinge, die er als Individuen oder als Klassenelemente erkennen könnte, eine potentiell unendliche Menge.

Die Theorie der semantischen Strukturen setzt die Existenz einer endlichen Menge von Namen und gleichfalls einer endlichen Menge von semantischen Merkmalen voraus. Wir wollen zeigen, daß zum induktiven Erkennen nur endlich viele Namen und semantische Merkmale erforderlich sind, obwohl es unbegrenzt viele Objekte zu identifizieren erlaubt. Zur Demonstration dient wieder das Zwanzig-Fragen-Spiel. Der Spieler errät ein Ding, das sich die anderen denken, mit höchstens zwanzig Fragen, auf die ihm die anderen mit Ja oder Nein antworten. Durch die Bejahung oder die Verneinung einer Frage wird eigentlich ein semantisches Merkmal des gesuchten Dinges bestimmt. Das

gesuchte Ding wird also durch höchstens zwanzig semantische Merkmale bestimmt. Und mit zwanzig Fragen kann es aus einer Menge von höchstens 2^{20} Elementen, doch nicht aus einer unendlichen Menge herausgefunden werden. Auch beim induktiven Erkennen bleibt also die Menge der Namen und die Menge der semantischen Merkmale endlich.

Der als Induktion aufgefaßte Erkennungsprozeß ist in seinem Wesen ein Erzeugungsprozeß. Der Gegenstand, den der Spieler des Zwanzig-Fragen-Spiels erraten soll, kann ein lebendes oder ein nichtlebendes Objekt, ein mythisches Wesen oder eine abstrakte Qualität sein. Wie auch immer, der Spieler kann seine Aufgabe erfolgreich lösen, wenn er über die zu erratende Sache etwas weiß, wenn er sich ihren Begriff bereits angeeignet hat. Wenn er aber über diesen Begriff nicht verfügte, könnte er die gesuchte Sache nicht finden, auch wenn es ihm gelänge, durch Fragen ihre wesentlichen Eigenschaften herauszubekommen. Die gewonnene Merkmalsliste hätte vom Standpunkt des Fragenden aus keinen Sinn. Aus entsprechenden Gründen kann man eine Sache nur dann erraten, wenn man auch mit den Eigenschaften, die man zu ihrer Bestimmung braucht, vertraut ist, d.h. wenn man sich auch *deren* Begriffe bereits angeeignet hat.

In der Theorie der semantischen Strukturen haben wir die Konzepte der Dinge als Namen und die Konzepte der Eigenschaften als semantische Merkmale bezeichnet. Was die Namen und Merkmale betrifft, so gibt es von ihnen, wie wir gezeigt haben, nur eine endliche Menge.

Wir haben also die Sachlage, daß ein Individuum mit einer begrenzten Menge an Vorwissen einem offenen Universum von potentiell identifizierbaren Objekten gegenübersteht. Aus diesem offenen Universum kann das Individuum nur eine endliche Anzahl von Objekten identifizieren. In einer nicht spezifizierbaren Zahl von Fällen wird es nicht zur Identifikation kommen. Aber auch die Zahl der möglichen erfolgreichen Identifikationen liegt nicht fest, denn in jeder konkreten Situation ist nur ein Ausschnitt aus dem offenen Universum relevant, und dieser Ausschnitt wird erst dadurch festgelegt, daß das Individuum die aktuelle Identifikationsaufgabe interpre-

tiert. Dies ist ein Argument dafür, daß man sagen kann, das induktive Erkennen sei ein kreativer Prozeß, ein Erzeugungsprozeß. Ein zweiter Aspekt der Kreativität des induktiven Erkennens besteht darin, daß das Subjekt - in Abhängigkeit von seinem Vorwissen und seiner Situationsinterpretation - Gegenstandseigenschaften des vorliegenden Objekts selektiert, identifiziert und zu einer "Gestalt" integriert, der beim Gelingen der Identifikation ein Name entspricht. Wir können also sagen, daß beim induktiven Erkennen das Gegenstandsmuster erst erzeugt wird. Zu demselben Schluß sind wir schon bei der Analyse der Semiotik des Erkennens im 3. Kapitel gekommen. Die Beziehung zwischen dem Gegenstand und seinem Muster wird durch das Gegenstandskonzept, den Namen, vermittelt. Der Name hat im Erkennungsprozeß die Funktion, die Selektion von Charakteristiken des Musters und deren Umkodierung in semantische Merkmale zu steuern.

Eine semantische Struktur ist keine feste, ein für allemal gegebene Organisation des die Identifikation gewährleistenden Gedächtnisses. Sonst hätte es gar keinen Sinn, induktive Erkennungsprozesse anzunehmen, man käme völlig mit den weniger anspruchsvollen deduktiven Verfahren aus. Ja, bei fester Struktur wäre ein induktives Verfahren überhaupt nicht möglich. Mit "semantischer Struktur" ist auch nicht die Anordnung der Gesamtheit aller im Wiedererkennungsgedächtnis gespeicherten Konzepte gemeint, sondern die durch die jeweils gestellte Identifikationsaufgabe determinierte Organisation einer Teilmenge der Namenmenge. Daß die semantische Struktur ein hierarchischer Zusammenhang semantischer Merkmale sei, wie vorhin gesagt wurde, ist nur ein anderer Ausdruck für denselben Sachverhalt.

Um zu zeigen, wie sich die semantische Struktur an der Identifikation beteiligt, müssen wir ein geeignetes Modell des Erkennungsprozesses aufstellen. Wir versuchen in diesem Modell, die hier diskutierten Aspekte des Erkennens mit den Mitteln der logischen Semantik auszudrücken. In erster Linie wollen wir die oben angedeuteten Prinzipien illustrieren, deshalb trachten wir nicht nach einem vollständigen logisch-semantischen Erkennungsmodell. Wir skizzieren vielmehr nur,

wie ein solches Modell konstruiert werden könnte. Ebendeswegen gehen wir auch von einer stark vereinfachten Ansicht der physiologischen Prozesse an den Rezeptoren aus.

Reize kommen vom Objekt an und verändern den Zustand der Rezeptorzellen. Die Veränderungen werden registriert und weiter verarbeitet. Die während der Darbietung des Objekts aufgenommene Information kann man sich vorstellen als Summe von Nachrichten des Typs: "Zelle z erregt", "Zelle z nicht erregt", "Gleichzeitig mit Zelle z wurde Zelle z' erregt", "Zellen eines bestimmten Rezeptorbereichs wurden gleichzeitig erregt". Die Rezeptorzellen registrieren die Reizintensität und nach manchen neurophysiologischen Befunden auch gewisse Beziehungen zwischen den stimulierten Zellen. Also auch diese Informationen müssen in dem vom Rezeptor abgegebenen Erregungsmuster enthalten sein.

Die Summe der Nachrichten, aus denen sich in unserem vereinfachten Modell die vom Rezeptor an die zentrale Verarbeitungsinstanz abgegebene Information aufbaut, können wir als ein Ensemble von Formeln der Prädikatenlogik notieren. Die Syntax des Erkennens suchen wir in den Mechanismen, die die Wohlgeformtheit dieser Formeln gewährleisten. Die Semantik des Erkennens suchen wir dagegen in Prozessen, die das Ensemble von Formeln in einen ganzheitlichen Ausdruck überführen, der erfüllbar ist. Der Erfüllbarkeit entspricht in unserem Fall die Übereinstimmung des ganzheitlichen Ausdrucks mit einem Namen. Nun kann das induktive Erkennen nach dem Vorbild des 20-Fragen-Spiels in der folgenden - hier nur angedeuteten - Richtung formalisiert werden.

Aus den Sinnesorganen treten in das Modell Teilinformationen über den zu identifizierenden Gegenstand ein. Diese Teilinformationen werden im Modell durch Formeln des Prädikatenkalküls erster Stufe dargestellt. Sie bilden zusammen die Ausgangsgesamtheit von Formeln. Ein beliebiges Element dieser Gesamtheit werde allgemein mit f bezeichnet. Die Gesamtheit der Formeln f - eine ungeordnete Aufzeichnung der aus der anliegenden sensorischen Stimulation kommenden Teilinformationen - ist das Material, aus dem der Name des Gegenstands zu bilden ist. Ein Name ist in diesem Modell eine komplexe Formel des Prädikaten-

kalküls erster Stufe. Der Name des Gegenstands wird schrittweise aufgebaut durch Vergleichung der sensorischen Teilinformationen mit semantischen Merkmalen.

Ein beliebiges semantisches Merkmal werde allgemein mit r bezeichnet. Ein semantisches Merkmal muß nicht einfach, d.h. durch einen Atomsatz darstellbar sein. r ist im allgemeinen eine komplexe erfüllbare Formel. Eine Formel heißt erfüllbar (im Sinne der von CARNAP herkommenden Semantiktradition), wenn sie wohlgeformt und in einer "möglichen Welt" gültig ist. In unserem Modell ist die Menge der möglichen Welten gleich der Menge der Namen, und die Entsprechung zwischen einer Formel und einem oder mehreren Namen gilt hier als Zuschreibung der Formel zu einer oder mehreren möglichen Welten. In CARNAP's Ausdrucksweise kann man eine solche Formel auch als Zustandsbeschreibung bezeichnen.

Die Hauptaufgabe unseres Modells ist es, die partiellen, durchweg einfachen Formeln f , welche die Sinnesdaten beschreiben, zu einer Formel anzuordnen, die mit einem der gespeicherten Namen identisch ist. Wenn das gelingt, dann ist das sensorische Muster identifiziert. Zu betonen ist hier noch einmal, daß die Überführung der vom Gegenstand ausgelösten Stimulation in ein komplexes ganzheitliches Gebilde durch die gespeicherten Namen vermittelt wird. Die Formeln f der Ausgangsgesamtheit können auf verschiedene Weise angeordnet werden. Um Eindeutigkeit zu garantieren, geht das Modell folgendermaßen vor:

Schritt 1: Es wird ein geeignetes, d.h. hinreichend allgemeines, semantisches Merkmal r ausgewählt. Diese Aktivierung eines semantischen Merkmals ist übrigens kein direktes Ergebnis der sensorischen Stimulation; denn semantische Merkmale sind zu verstehen als Konzepte von Gegenstandseigenschaften.

Schritt 2: Das gewählte semantische Merkmal r wird mit jeder Formel f der Ausgangsgesamtheit verglichen, d.h. für das gewählte r und alle f wird der Wahrheitswert der Konjunktion $f \cdot r$ bestimmt. Wenn alle Konjunktionen $f_1 \cdot r, f_2 \cdot r, f_3 \cdot r, \dots$ wahr sind, so wird dem zu identifizierenden Muster das semantische Merkmal r zugeschrieben. Wenn dagegen alle diese Konjunktionen falsch sind, so wird dem Muster das Merkmal r nicht

zugeschrieben. Und wenn keiner der beiden Fälle eintritt, sondern einige der Konjunktionen wahr, andere falsch sind, so ist das Merkmal r für das vorliegende Muster irrelevant. In diesem dritten Fall muß Schritt 1 wiederholt werden.

Schritt 3: Wenn schon mindestens ein semantisches Merkmal r des zu identifizierenden Musters gefunden ist, so wird ein weiteres Merkmal unter denen ausgewählt, die im Graphen einer hierarchischen Struktur aus den anstehenden Merkmalen Nachfolger von r sein können. Mit diesem neu gewählten Merkmal wird Schritt 2 wiederholt.

Schritt 4: Die mehrfache Wiederholung der Schritte 2 und 3 kann zur Bestimmung einer ganzen Merkmalsliste, also eines Namens für das Muster, führen, sie muß aber nicht immer mit einem Erfolg enden. Mit der Bestimmung eines Namens ist der Identifikationsprozeß beendet.

Zwei kurze Bemerkungen:

1. Dieses Modell stellt keinen rein syntaktischen Entscheidungsprozeß dar. Vielmehr erfordert die Transformation der Ausgangsgesamtheit von Formeln f in eine komplexe erfüllbare Formel auch semantische Mittel, nämlich die Prüfung der Erfüllbarkeit der konstruierten Formeln.

2. Die Ausgangsdaten werden wirklich in eine ganzheitliche Anordnung transformiert, d.h. es wird nicht bloß eine Konjunktion oder Liste von Teildaten aufgebaut.

Die soeben formulierten Prinzipien zeigen wohl andeutungsweise, wie eine hierarchische semantische Struktur die Identifikation der Muster aller derjenigen Objekte gewährleisten kann, deren Konzepte (Namen) sich das erkennende Subjekt bereits angeeignet hat. Voraussetzung ist allerdings, daß die semantischen Merkmale die im 3. und 6. Kapitel aufgestellten Kriterien der Hierarchizität erfüllen. Dann können die semantischen Merkmale durch erfüllbare Formeln dargestellt werden (wobei Erfüllbarkeit heißt, daß es eine Teilmenge von Namen gibt, deren Elemente alle dieses Merkmal tragen). Die Menge der Namen in einer semantischen Struktur stellt eigentlich die Menge der möglichen Resultate von Erkennungsprozessen dar. Sie entspricht also dem Konzept der "Menge der möglichen

Welten". Deshalb können wir die Erfüllung definieren als die Entsprechung mit einem Namen oder einer Teilmenge von Namen.

5.3 Semantische Merkmale

Semantische Merkmale sind die Grundbausteine der semantischen Struktur. Wir wollen hier zusammenfassen, was wir bisher über ihre Beschaffenheit gesagt haben. Ein semantisches Merkmal ist das Konzept einer Gegenstandseigenschaft. Die Wahrnehmungsgegenstände haben physikalische, chemische und andere Eigenschaften. Bei der sensorischen Abbildung des Gegenstands auf den Rezeptoren werden seine Eigenschaften durch Charakteristika des sensorischen Reizmusters widerspiegelt. Bei der Identifikation wird das primäre sensorische Gegenstandsabbild, das eine partikuläre und in der Regel unvollständige Kollektion von Charakteristika bildet, mit Hilfe der semantischen Struktur in eine Kollektion semantischer Merkmale transformiert. Da an der Transformation von Charakteristika in Merkmale das Konzept des betreffenden Gegenstands beteiligt ist, garantiert diese Transformation, daß Gegenstand und Muster sich entsprechen. Die Entsprechung wird durch den Gegenstandsbegriff vermittelt. In ZINČENKO's (1973) Terminologie ausgedrückt, garantiert die Transformation der Charakteristika in Merkmale eine gegenseitige Angleichung des Objektfeldes und des Phänomenfeldes.

Die Beziehungen zwischen den semantischen Merkmalen bilden insgesamt die semantische Struktur. Da die semantische Struktur des Wiedererkennungsgedächtnisses nach den Ergebnissen unserer Experimente hierarchisch ist, stehen die semantischen Merkmale in Beziehung der Über- und Unterordnung. Auf einer semantischen Merkmalsliste, die einen Namen (den Begriff eines individuellen Gegenstands) darstellt - beim Identifizieren ist sie das Ergebnis der Transformation eines sensorischen Reizmusters - stehen die semantischen Merkmale miteinander in einem hierarchischen Zusammenhang. Allerdings brauchen wir, wenn wir das Erkennen als einen induktiven Pro-

zeß auffassen, offenbar nicht zu fordern, daß alle im Wiedererkennungsgedächtnis gespeicherten semantischen Merkmale und Namen eine strikte Hierarchie bilden. Die Menge der Merkmale zerfällt in Ebenen, auf denen jeweils hierarchische Anordnung gilt. Wenn der induktive Identifikationsprozeß ebenenweise verläuft, dann wird jeweils eine als semantische Struktur geordnete Teilmenge der Merkmale aktualisiert.

Wir haben die semantischen Merkmale durch ihre Funktion im Erkennungsprozeß definiert. Es resultiert ein ziemlich abstraktes Bild von ihnen, das über ihre psychologische Natur wenig aussagt. Eine nähere Kenntnis der psychologischen Natur der semantischen Merkmale dürfte für die Vertiefung der Einsicht in die Semiotik des menschlichen Verhaltens wichtig sein. Einschlägige Erkenntnisse können wir z.B. mit der Methode der Gruppierung von Wörtern nach Bedeutungsähnlichkeit, die wir zur Messung semantischer Distanzen benutzt haben, gewinnen. Jeder Knoten des hierarchischen Graphen, der die Gruppierungsergebnisse darstellt, repräsentiert ein semantisches Merkmal. Ähnliche Einsichten gibt uns auch das Zwanzig-Fragen-Spiel. KINTSCH (1970) hat gezeigt, daß die Identifikation eines gesuchten Dinges in diesem Spiel über dieselben semantischen Merkmale erfolgt wie das Gruppieren von Wörtern nach Bedeutungsähnlichkeit.

Eine weitere Möglichkeit, Kenntnisse über die psychologische Natur der semantischen Merkmale zu gewinnen, bietet sich in der Verwendung unbestimmter Reizsituationen, wie Rorschach-Tafeln oder unvollendeter Sätze, an. Die Reaktion von Versuchspersonen auf unbestimmte, mehrdeutige Reizsituationen ermöglicht es zu untersuchen, wie Reizcharakteristika in semantische Merkmale transformiert werden. Auf diesem Prinzip beruht die Rorschach-Diagnostik.

5.4 Möglichkeiten und Grenzen der hierarchischen Struktur

Trotz aller Belege, die wir sammeln konnten, dürfen wir die Hierarchizität der semantischen Struktur wohl nicht für das

Universalprinzip der Gedächtnis- und Erkennungsprozesse halten. G.A. MILLER (in: MILLER & MCNEILL 1969) hat z.B. seine Methode der semantischen Distanzmessung außer auf Substantive auch auf andere Wortarten angewandt. Während sich Adjektive ebenso leicht wie Substantive hierarchisch skalieren ließen, versagte die Methode bei Verben. Von ähnlicher Art scheint ein Resultat zu sein, zu dem mich mein Interesse am Zwanzig-Fragen-Spiel gebracht hat. Nach bisher unvollständigen Befunden ist es anscheinend praktisch unmöglich, die Spielaufgabe zu lösen, wenn statt eines Dinges ein Ereignis erraten werden soll. Das würde bedeuten, daß das induktive Erkennen ebenso beschränkt ist wie die hierarchische Speicherung von Wissen: Weder das eine noch das andere bezieht sich auf Ereignisse.

Den Prozeß der wechselseitigen Beeinflussung von Objekt und Konzept, den das induktive Erkennen voraussetzt, versuchte ich an der Formwahrnehmung bei Rorschach-Tafeln zu studieren. Dem Experiment unterzogen sich zwei professionelle bildende Künstler, Männer im Alter von 30 Jahren. Der Grund, gerade Künstler um die Teilnahme am Versuch zu bitten, war, daß das Verfahren von der Versuchsperson eine gewisse zeichnerische Geschicklichkeit verlangte.

Bei der Rorschach-Untersuchung wird der Proband mitunter aufgefordert, seine Perzepte in eine Fotokopie der Tafel einzuzichnen. Das dient zur exakteren Bestimmung der Beschaffenheit der Perzepte. In unserem Experiment wurden den Probanden nach Abschluß der Untersuchung die Tafeln erneut vorgelegt. Nun sollten sie auf jeder Tafel ein Perzept festlegen, eventuell eine der Deutungen, die sie schon während der Untersuchung gegeben hatten. Sie sollten es in der Fotokopie der Tafel einzeichnen und außerdem noch so, wie sie es auf der Tafel sahen, auf ein Blatt Papier zeichnen. Diese Aufgabe war so anspruchsvoll und langwierig, daß nur zu wenig mehr als der Hälfte der Tafeln die Perzepte vollständig ausgearbeitet werden konnten.

Die Mehrzahl der Zeichnungen enthielt Details, die man nur schwer in der Struktur des gedeuteten Kleckses identifizieren konnte oder die in der Struktur des Kleckses überhaupt keine

Entsprechung hatten. Das Konzept hat, so scheint es, das Wahrnehmungsobjekt zu einer regelmäßigeren und konkreteren Gestalt vervollkommenet. Ein extremes Beispiel hierfür ist die Ausarbeitung des linken oberen Details von Tafel IV (Lage der Tafel: rechte Seite der Ausgangslage nach oben gedreht). Proband K. nahm dies Detail als einen Mövenkopf wahr (einschließlich der weißen Fläche innen). Beim Einzeichnen in die Fotokopie des Kleckses zog er den Umriß des Perzepts und betonte einige Details am Umriß. Beim nachfolgenden Zeichnen arbeitete er den Vogelkopf detailliert und der Struktur des Kleckses tatsächlich entsprechend aus, aber dann ergänzte er ihn plötzlich mit einem etwas schematisierten Körper.

Bei den Perzepten, zu denen keine Zeichnung zu bekommen war, betonten die Versuchsteilnehmer meistens, daß die Eintragung in die Fotokopie alles Wesentliche des Perzepts enthalte und daß es überflüssig sei, das Perzept noch einmal gesondert zu zeichnen. In einigen Fällen konnte man die Erfassung des Perzepts in der Fotokopie wirklich für eine grobe Zeichnung halten. In zwei, drei Fällen aber war eine solche Behauptung nicht zu akzeptieren. Als Beispiel führen wir diesmal die Bearbeitung von Tafel I (Ganzauffassung, in Ausgangslage) durch den Probanden J. an. J. faßte die Tafel als "zwei Tanzpaare, kleine Frauen in der Mitte" auf. Das ist eine gute, kreative Bearbeitung des Kleckses. Auf der Fotokopie dagegen zog der Proband nur den Rand des Kleckses aus und weigerte sich, mit der oben angeführten Begründung, weiter zu zeichnen. Im Unterschied zum vorigen Fall, wo das Wahrnehmungsobjekt vom Konzept her ergänzt wurde, hat sich hier anscheinend das Konzept unter dem Einfluß des Objekts aufgelöst und verwischt.

Die Deutung von Rorschach-Klecksen ist eine schwierige Wahrnehmungsaufgabe. Das Wahrnehmungsobjekt kann keinem der Konzepte, die wir zum Erkennen einsetzen, voll entsprechen; denn es ist selbst unbestimmt. Wenn das Erkennen als ein gegenseitiges Sichangleichen von Konzept und sensorischer Stimulation verläuft, wie wir es in dieser Arbeit erörtert haben, dann kann es zur Identifikation einer Form in einem

Rorschachschen Tintenkleecks im wesentlichen auf zwei Arten kommen. Entweder ergänzt das Konzept das sensorische Reizmuster um Elemente, die in den sensorischen Daten selbst nicht enthalten sind, oder, umgekehrt, die sensorischen Daten erzwingen eine Auflösung und Verwischung des Konzepts derart, daß das unvollständige, fragmentarische Reizmuster schließlich akzeptabel wird. Daß beide Lösungen der Identifikationsaufgabe existieren, haben wir mit diesem letzten Versuch belegt. Auch dies spricht für die Annahme, daß das Gedächtnis hierarchische Strukturen enthält, die beim induktiven Erkennen eingesetzt werden.

Zwar ist dieser kleine Versuch kein zwingender Nachweis hierarchischer Strukturen im Gedächtnis; wohl aber stützt er die Auffassung, daß die Bildung eines Perzepts aus der Wechselwirkung von Gegenstand, Reizmuster und Konzept resultiert, wobei die physische Stimulation und die konzeptuelle Basis durchaus im Wettstreit um ein Perzept liegen können.

6. EINE FORMALE THEORIE DER SEMANTISCHEN STRUKTUREN

Die im 3. Kapitel eingeführte Definition der semantischen Strukturen beruht auf gewissen Beziehungen zwischen einer Menge von Namen, einer Menge von semantischen Merkmalen und einem semantischen Distanzmaß. Diese Beziehungen sind in der natürlichen Sprache schwer exakt zu beschreiben. Deshalb wollen wir die Theorie der semantischen Strukturen in diesem Kapitel als eine formale Theorie einführen. Die Formalisierung der Theorie der semantischen Strukturen ermöglicht eine präzise Definition der hierarchischen und der paradigmatischen Struktur, aber sie hat außerdem noch zwei weitere Ziele. Erstens liefert sie einen Existenzbeweis für Strukturen mit solchen Eigenschaften, zweitens erlaubt sie die Ableitung einiger weiterer Eigenschaften von semantischen Strukturen, bezüglich derer (der Eigenschaften) wir die im 4. Kapitel dargestellten Experimente interpretiert haben.

6.1 Ein Differenzmaß

Semantische Strukturen hängen eng mit einem semantischen Distanzmaß zusammen. Ein solches Maß hat einige spezifische Eigenschaften. Ein Maß mit diesen Eigenschaften nennen wir "Differenzmaß" und führen es allgemein als Maß auf einer nicht-leeren endlichen Menge ein.

Definition 1: Gegeben sei eine endliche Menge $M \neq \emptyset$, ein Teilmengensystem $\mu, \mu = \{ M_i / M_i \subseteq M, i=1, \dots, n \}$ sowie eine Metrik ρ auf $M \times M$. Die Metrik ρ nennen wir "zum System μ gehöriges Differenzmaß auf M " genau dann, wenn für zwei beliebige Elemente a, b aus M $\rho(a,b) = f(n(a,b))$ gilt, mit $n(a,b) = \text{card} \{ M_i \subset \mu / a \in M_i \text{ \& } b \in M_i \}$.

Satz 1: Für ein beliebiges Teilmengensystem μ der endlichen nichtleeren Menge M existiert mindestens ein zu μ gehöriges Differenzmaß auf M .

Beweis: Bezeichnen wir mit N die maximale Mächtigkeit derjenigen Teilmenge aus μ , deren Durchschnitt mit jeder anderen Teilmenge nichtleer ist. (Für alle a, b aus M gilt $n(a,b) \leq N$).

1. Wenn μ ein Teilmengensystem mit $N = 0$ ist (d.h. wenn μ entweder leer ist oder die leere Menge enthält), dann genügt es, $\rho(a,b) = c$, für $a \neq b$, zu setzen, wo c eine beliebige Konstante ist und $\rho(a,a) = 0$. Dann ist ρ offensichtlich eine Metrik. Um die Definitionsbedingung zu erfüllen, braucht nur eine geeignete Funktion f mit den geforderten Eigenschaften gefunden zu werden, für die $f(0) = c$ gilt (denn für $a = b$ ist $n(a,b)=0$).

2. Sei $N \neq 0$. Wir definieren

$$\rho(a,b) = (2N - n(a,b)) / N, \quad \rho(a,a) = 0.$$

Angesichts dessen, daß $n(a,b) \leq N$ ist, gilt $\rho(a,b) \geq 1$. Also impliziert $\rho(a,b) = 0$, daß $a = b$ ist. Aus der Definition von $n(a,b)$ wird auch evident, daß ρ symmetrisch ist.

Wir wählen drei Elemente $a, b, c \in M$ und zeigen $\rho(a,c) \leq \rho(a,b) + \rho(b,c)$. Wenn zwei Elemente identisch sind, gilt diese Ungleichung. Es seien also a, b, c paarweise verschieden. Offensichtlich gilt

$$2N + n(a,c) \geq n(a,b) + n(b,c) \quad \text{und daraus folgt}$$

$$\frac{2N - n(a,c)}{N} \leq \frac{2N - n(a,b)}{N} + \frac{2N - n(b,c)}{N}$$

Daher ist ρ eine Metrik, und da $f(x) = \frac{2N - x}{N}$, ist f eine monoton fallende nichtnegative Funktion.

6.2 Semantische Strukturen

Bezeichnen wir die Menge der Namen mit J und die Menge der Merkmale mit Φ . $J = \{i, j, k, \dots\}$ und $\Phi = \{\varphi, \chi, \psi, \dots\}$ seien endliche Mengen. Ferner bezeichnen wir mit ρ die Metrik und mit F ein Teilmengensystem der Menge J , $F = \{J_i / i=1, \dots, n\}$.

Definition 2: Eine semantische Struktur ist ein Tripel $\langle J, \Phi, \rho \rangle$, für welches folgendes gilt:

(1) Das Teilmengensystem F läßt sich auf die Menge Φ eindeutig abbilden. Für $\varphi \in \Phi$ und ein $J_i \subseteq J$, das auf dieses φ abgebildet ist, schreiben wir $J_i = \varphi(J)$. Für alle $\varphi \in \Phi$ gilt $\varphi(J) \neq \emptyset$. Wir sagen, alle Namen i aus $\varphi(J)$ haben das Merkmal φ . Weiter bezeichnen wir mit γ das Merkmal, für welches $\gamma(J) = J$ gilt.

(2) Die Metrik ρ ist ein Differenzmaß auf J , das zum System F der Teilmengen von J gehört.

Wir definieren nun eine hierarchische Struktur mit denselben Kriterien wie in Kapitel 2.

Definition 3: Eine hierarchische Struktur ist eine semantische Struktur, in der zu jedem $\varphi(J)$ eine nichtnegative reelle Zahl δ existiert derart, daß für jedes Paar $i, j \in \varphi(J)$ $\rho(i, j) \leq \delta$ und außerdem für jedes Paar i, k , mit $i \in \varphi(J)$ und $k \notin \varphi(J)$, $\rho(i, k) > \delta$ gilt.

Aus Definition 3 leiten wir zwei Sätze ab, die die beiden Haupteigenschaften einer hierarchischen Struktur beschreiben.

Satz 2: Eine semantische Struktur ist eine hierarchische Struktur genau dann, wenn für jedes Paar von Teilmengen $\varphi(J), \chi(J) \subseteq J$ gilt, daß entweder ihr Durchschnitt leer ist oder daß die eine eine Teilmenge der anderen ist.

Beweis: Schreiben wir $\varphi(J) = J_i$, $\chi(J) = J_j$.

A. Es gebe in der hierarchischen Struktur mindestens einen Fall, in dem der Durchschnitt $J_i \cap J_j$ nichtleer ist, während J_i, J_j unterschiedlich sind und keine von beiden eine Teilmenge der anderen ist. Nach Definition 3 gilt dann: Zur Teilmenge J_i existiert eine Zahl δ' derart, daß für zwei beliebige Elemente $i, j \in J_i$ das Maß $\rho(i, j) \leq \delta'$ ist, und zur Teilmenge J_j gibt es eine Zahl δ'' derart, daß für zwei beliebige Elemente $k, l \in J_j$ das Maß $\rho(k, l) \leq \delta''$ ist. Sei nun h ein Element des Durchschnitts $J_i \cap J_j$. Dann gilt $\rho(i, h) \leq \delta'$ und $\rho(k, h) \leq \delta''$.

(1) Wenn $\delta' = \delta''$, dann ist $J_i = J_j$ und daher $J_i \subseteq J_j$.

(2) Wenn $\delta' < \delta''$, dann ist $\rho(i, h) < \delta''$, denn $\rho(i, h) \leq \delta' < \delta''$. Daraus folgt, daß alle Elemente der Teilmenge J_i Elemente der Teilmenge J_j sind, d.h. $J_i \subseteq J_j$ und daher $J_i \subseteq J_j$.

(3) Wenn $\delta' > \delta''$, dann führen wir den Beweis analog wie in (2), und offensichtlich ist dann $J_j \subseteq J_i$. Die Annahme hat also zum Widerspruch geführt.

B. Für jedes Paar von Teilmengen $J_i, J_j \subseteq J$ sei entweder ihr Durchschnitt leer, oder die eine sei Teilmenge der anderen. Dann gibt es nach der Definition 3 zu J_i eine Zahl δ' und zu J_j eine Zahl δ'' derart, daß für jedes Paar i, j aus J_i und jedes $h \notin J_i$ gilt $\rho(i, j) \leq \delta'$ und $\rho(i, h) > \delta'$; ferner für jedes Paar k, l aus J_j und jedes $h \notin J_j$ $\rho(k, l) \leq \delta''$ und $\rho(k, h) > \delta''$. Denn ohne Einschränkung der Allgemeinheit haben die i aus J_i alle diejenigen Merkmale φ' gemeinsam, für die $J_i \subseteq \varphi'(J)$ gilt, während die h , die nicht Elemente von J_i sind, mit den Elementen von J_i höchstens die Merkmale φ'' gemeinsam haben, für welche gilt $J_i \subset \varphi''(J)$. Daraus folgt, daß die semantische Struktur hierarchisch ist.

Satz 3: Das Maß ρ in der hierarchischen Struktur erfüllt die ultrametrische Ungleichung

$\rho(i, k) \leq \max [\rho(i, j), \rho(j, k)]$ für beliebige Tripel $i, j, k \in J$.

Beweis: Wenn das Maß ρ die ultrametrische Ungleichung erfüllen soll, dann müssen offensichtlich die folgenden drei Ungleichungen gleichzeitig erfüllt sein:

$$\rho(i,k) \leq \max [\rho(i,j), \rho(j,k)]$$

$$\rho(i,j) \leq \max [\rho(i,k), \rho(j,k)]$$

$$\rho(i,k) \leq \max [\rho(i,j), \rho(i,k)] .$$

Damit diese Ungleichungen gleichzeitig erfüllt sind, ist es notwendig und hinreichend, daß zwei der Distanzen zwischen den Namen i , j und k gleich sind und die dritte kleiner als die beiden ist.

Nehmen wir an, das Maß ρ in der hierarchischen Struktur erfülle die ultrametrische Ungleichung nicht. Dann gilt, ohne Einschränkung der Allgemeinheit,

$\rho(i,k) < \rho(j,k)$ und $\rho(i,k) < \rho(i,j)$. Aus diesen beiden Ungleichungen und Definition 2 folgt, daß es ein Merkmal φ mit der Eigenschaft $i, j \in \varphi(J)$ und $k \notin \varphi(J)$ gibt, sowie ein Merkmal χ mit der Eigenschaft $j, k \in \chi(J)$ und $i \notin \chi(J)$.

Der Durchschnitt $\varphi(J) \cap \chi(J)$ ist dann offenbar nicht leer, und keine der beiden Teilmengen ist eine Teilmenge der anderen - das ist ein Widerspruch.

Definition 4: Eine paradigmatische Struktur ist eine semantische Struktur, die nicht hierarchisch ist.

Lemma 1: In einer paradigmatischen Struktur existiert mindestens ein Merkmalspaar $\varphi, \chi \in \Phi$ mit nichtleerem Durchschnitt $\varphi(J) \cap \chi(J)$, bei dem weder $\varphi(J)$ eine Teilmenge von $\chi(J)$ noch $\chi(J)$ eine Teilmenge von $\varphi(J)$ ist.

Beweis: Lemma 1 ist die Negation von Satz 2. Daraus und aus Definition 4 folgt die Gültigkeit des Lemmas.

Die Umkehrung zu Satz 3 gilt nicht. Ein Maß in einer paradigmatischen Struktur kann also ultrametrisch sein.

6.3 Der Graph einer semantischen Struktur

Bisher haben wir semantische Strukturen definiert und eini-

ge ihrer Eigenschaften abgeleitet, doch noch keinen Existenzbeweis geführt. Zum Beweis der Existenz semantischer Strukturen genügt es, die Existenz ihrer Modelle zu zeigen. Solche Modelle sind beispielsweise die Graphen semantischer Strukturen.

Sei F ein Teilmengensystem der Menge J , mit der Eigenschaft der einfachen Abbildbarkeit auf die Menge Φ der Merkmale. Sei ferner Γ eine Abbildung von F , mit der Eigenschaft, daß für jedes Paar ungleicher Elemente J_i, J_j des Systems F genau dann $J_i \in \Gamma J_j$ gilt, wenn der Durchschnitt $J_i \cap J_j$ nichtleer ist.

Definition 5: Der Graph einer semantischen Struktur ist ein Paar $\langle F, \Gamma \rangle$, wo F die Menge der Knoten des Graphen und Γ eine eindeutige Abbildung dieser Menge in sich bedeutet.

Eine Kante des Graphen ist die Verbindung zwischen zwei Knoten. Zwei Knoten $J_i, J_j \in F$ werden genau dann durch eine Kante verbunden, wenn $J_i \in \Gamma J_j$ ist. Aus der Beschaffenheit der Abbildung Γ folgt offensichtlich, daß $J_i \in \Gamma J_j$ genau dann gilt, wenn $J_j \in \Gamma J_i$ gilt. Ein Baum ist ein Graph, in dem jedes Knotenpaar durch genau einen Weg verbunden ist, wobei als Weg eine zusammenhängende Folge von Kanten des Graphen gilt. Wir wollen zeigen, daß der Graph einer hierarchischen Struktur ein Baum ist.

Satz 4: Der Graph einer hierarchischer Struktur ist ein Baum, und der Graph einer paradigmatischen Struktur ist kein Baum.

Der Beweis ergibt sich direkt aus Satz 2, Lemma 1 und der Existenz des Merkmals γ in der semantischen Struktur. Die Existenz des Merkmals γ garantiert den Zusammenhang (connectedness) des Graphen der semantischen Struktur, und Satz 2 garantiert, daß der Graph einer hierarchischen Struktur keine Zyklen enthält. Lemma 1 dagegen garantiert, daß der Graph einer paradigmatischen Struktur mindestens einen Zyklus enthält.

Mit Hilfe des Graphen der semantischen Struktur können wir eine Reihe weiterer Begriffe einführen. Den Knoten J_i nennen wir genau dann einen Vorgänger des Knoten J_j , wenn $J_j \in \Gamma J_i$ und $J_j \subseteq J_i$ ist. J_j nennen wir in diesem Fall einen Nachfolger

von J_i . Die Namen selbst, d.h. die einelementigen Teilmengen der Menge J , haben keine Nachfolger. Den Knoten ohne Vorgänger nennen wir die Quelle des Graphen. Das ist in einer hierarchischen Struktur offensichtlich der Knoten $\gamma(J)$.

Ferner nennen wir "unmittelbare Nachfolger" des Knoten J_i diejenigen unter dem Knoten J_j , die folgende Eigenschaften haben:

- a) sie sind Nachfolger von J_i ,
- b) es gibt keinen Knoten J_k derart, daß J_k Nachfolger von J_i und J_j Nachfolger von J_k wäre.

Im Graphen einer hierarchischen Struktur kann man die Knoten mittels der Relation "Vorgänger - Nachfolger" ordnen. Sei λ eine Funktion, die jedem Knoten auf die folgende Weise eine positive Zahl zuordnet: Die Quelle bekommt die Zahl 1. Wenn y der unmittelbare Nachfolger von x ist, so ist $\lambda(y) = \lambda(x) + 1$. Die Knoten J_i und J_j , für die $\lambda(J_i) = \lambda(J_j)$ gilt, bezeichnen wir als Knoten desselben Niveaus.

Übersetzt aus dem Tschechischen von W. Matthäus

Literatur

- BLEDSE, W.W., BROWNING, I., Pattern recognition and reading by machine. In: Proceedings of the Eastern Joint Computer Conference 1959, 16, 225-232
- BRUNER, J.S., GOODNOW, J.J., AUSTIN, G.A., A study of thinking. New York, Wiley 1967 (1. Auflage 1956)
- CLIFF, N., Scaling. In: Annual Review of Psychology 1973, 24, 473-506
- CORCORAN, D.W.J., Serial and parallel classification. British Journal of Psychology 58, 1967, 197-203
- CORCORAN, D.W.J., Pattern Recognition. Harmondsworth, Penguin 1971
- EGETH, H.E., Parallel versus serial processes in multidimensional stimulus discrimination. Perception and Psychophysics 1, 1966, 245-252
- FILLENBAUM, S., Psycholinguistics. Annual Review of Psychology 22, 1971, 251-308
- HALMIOVA, O., ŠIPOŠ, I., Intentional forgetting of categorized words. Studia Psychologica 18, 1976, 183-190
- HICK, W.E., On the rate of gain of information. Quarterly Journal of Experimental Psychology 4, 1952, 11-26
- HUNT, E.B., Concept Learning: An information processing problem. New York, Wiley 1962
- HUNT, E.B., MARIN, J., STONE, P.I., Experiments in induction. New York, Wiley 1966

- JOHNSON, S.C., Hierarchical clustering schemes. *Psychometrika* 32, 1967, 241-254
- KINTSCH, W., Effects of memory structure upon free recall, word identification and classification tasks. Proceedings of the International Conference on Human Learning. Institute of Psychology, Czechoslovak Academy of Sciences, Prague 1970, Vol. II, 215-224
- KOLMAN, L., Normative approach to perceptual learning. Proceedings of the International Conference on Human Learning. Institute of Psychology, Czechoslovak Academy of Sciences, Prague 1970, 267-272
- KOLMAN, L., Toward the structural basis of cognition. Proceedings of the Second Prague Conference on Human Learning and Problem Solving. Universita Karlova, Prague 1974, 207-209
- KOLMAN, L., An attempt to investigate human memory structure. A paper held at the Third Prague Conference on the Psychology of Human Learning and Problem Solving. Prague 1977
- KÜLLER, R., A semantic model for describing perceived environment. Document D 12: 1972 National Swedish Building Research. Stockholm 1972
- LIBERMAN, A.M., COOPER, F.S., HARRIS, K.S., MacNEILAGE, P.F., STUDDERT-KENNEDY, M., Some observations on a model for speech perception. In: WATHEN-DUNN, W. (Ed.), *Models for the Perception of Speech and Visual Form*. Cambridge, Mass., MIT Press 1964, 68-87
- LIBERMAN, A.M., COOPER, F.S., SHANKWEILER, D.P., STUDDERT-KENNEDY, M., Perception of the speech code. *Psychological Review* 74, 1967, 431-461

- LIBERMAN, A.M., DELATTRE, P.C., COOPER, F.S., The rule of selected stimulus variables in the perception of the unvoiced stop consonants. *American Journal of Psychology* 65, 1952, 497-516
- LYNCH, K. *The image of the city*. Cambridge, Mass., MIT Press 1960
- MEADOW, C.T., *The analysis of information systems. A programmer's introduction to information retrieval*. New York, Wiley 1967
- MILLER, G.A., Psycholinguistic approaches to the study of communication. In: ARM, D.L. (Ed.), *Journeys in Science*. Albuquerque, The University of New Mexico Press 1967, 22-73
- MILLER, G.A., A psychological method to investigate verbal concepts. *Journal of Mathematical Psychology* 6, 1969, 169-191
- MILLER, G.A., McNEILL, D., Psycholinguistics. In: LINDZEY, G., ARONSON, E. (Eds.), *The Handbook of Social Psychology* Vol. 3, Reading, Mass., Addison-Wesley 1969, 666-794
- NEISSER, U., Decision time without reaction time: Experiments in visual scanning. *American Journal of Psychology* 76, 1963, 376-385
- NEISSER, U., NOVICK, R., LAZAR, R., Searching for ten targets simultaneously. *Perceptual and Motor Skills* 17, 1963, 955-961
- NICKERSON, R.S., FEEHRER, C.E., Stimulus categorization and response time. *Perceptual and Motor Skills* 18, 1964, 785-793
- OSGOOD, C.E., SUCI, G.J., TANNENBAUM, P.H., *The measurement of meaning*. Urbana, The University of Illinois Press 1957

RABBITT, P.M., Learning to ignore irrelevant information.
American Journal of Psychology 80, 1967, 1-13

ROSENBLATT, F., The perceptron: A probabilistic model for information storage and organization in the brain. Psychological Review 65, 1958, 386-408

STERNBERG, S., High speed scanning in human memory. Science 153, 1966, 652-654

STERNBERG, S., Two operations in character recognition: Some evidence from reaction time measurements. Perception and Psychophysics 2, 1967, 45-53

TOLMAN, E.C., Cognitive maps in rats and man. Psychological Review 55, 1948, 189-208

UHR, L., "Pattern recognition" computers as models for form perception. Psychological Bulletin 60, 1963, 40-73

UHR, L., Pattern recognition, learning, and thought. Englewood Cliffs, Prentice-Hall 1973

UNGER, S.H., Pattern recognition and detection. Proc. IRE 47, 1959, 1737-1752

UTTLEY, A.M., Conditional probability computing in a nervous system. In: Mechanisation of Thought Processes. London, Her Majesty's Stationary Office 1959, 119-147

UTTLEY, A.M., The transmission of information and the effect of local feedback in theoretical and neural networks. Brain Research 2, 1966, 21-50

WATANABE, S., Une explication mathématique du classement d'objets. In: DOCKX, S., BERNAYS, P. (Eds.), Information and Prediction in Science. New York, Academic Press 1965, 39-76

WELFORD, A.T., The measurement of sensory-motor performance: Survey and reappraisal of twelve years' progress. Ergonomics 3, 1960, 189-230

ZINČENKO, V.P., Productive perception. Československá Psychologie 17, 1973, 306-320

ZINČENKO, V.P., VERGILES, N.Ju., Formation of visual images. Studies of stabilized retinal images. New York, Consultants Bureau 1972

MATHEMATISIERUNGSTENDENZEN IN DER MUSIKWISSENSCHAFT:

EINE BIBLIOGRAPHIE ZUR QUANTITATIV-STATISTISCHEN UND ALGEBRAISCH-FORMALEN ANALYSE MUSIKALISCHER STRUKTUREN

T.H. Stoffer, Bochum

M.G. Boroda, Tbilisi

Musikwissenschaftler scheinen ein tief in der Historie verwurzeltetes Mißtrauen gegenüber quantitativ- und formal-analytischen Methoden der musikalischen Form- und Strukturanalyse zu besitzen. An der weltweit verschwindend kleinen Zahl von Autoren, die sich quantitativer oder algebraischer Analysemethoden bedienen, im Vergleich zur Zahl der Autoren, die sich der traditionellen, historisch-hermeneutischen Methode verpflichtet fühlen, erkennt man, daß die Verteilung dieses Mißtrauens globaler Art ist. So waren es dann auch überwiegend an der Musik interessierte Wissenschaftler anderer Disziplinen, die quantitative oder algebraische Beschreibungs- und Analysemethoden aus ihrer Disziplin erstmalig auf musikalische Strukturen zu übertragen versuchten.

Die Defizite musikwissenschaftlicher Analysemethoden werden auch von Seiten der Musikwissenschaft bedauert (WIORA, 1970), aber nach wie vor versteht sich die Musikwissenschaft als überwiegend zu den Geisteswissenschaften gehörig, was zwar die Anwendung quantitativer Methoden nicht auszuschließen brauche (WIORA, 1970), ihnen damit aber faktisch eine untergeordnete Rolle gegenüber der historisch-hermeneutischen Vorgehensweise zuweist, da sich die Reichweite ihrer Anwendungen auf Musik erst noch erweisen müsse (WIORA, 1970). Gemeint ist hier von WIORA in erster Linie die informationstheoretisch-statistische Analysemethoden, die bis heute die einzige in größerem Umfang angewendete quantitative Methode im Bereich der Musikwissenschaft darstellt (vgl. Abschnitt I. der Bibliographie). Ihre "Bewährungsprobe" hat sie inzwischen hinter sich, wobei die zu Tage geförderten Resultate u.E. nur die skeptische Einschätzung durch WIORA unterstützen können, da sich die informationstheoretische Analysemethoden als nicht deskriptiv adäquat erwiesen hat (vgl. COHEN, 1962 in der Bibliographie; de la MOTTE-HABER, 1972,

1976; STOFFER, 1979 in der Bibliographie).

Die erwiesene Inadäquatheit einer Methode sollte aber nicht zum Anlaß genommen werden, jegliche Art quantitativer oder algebraischer Methoden im Bereich der Musikwissenschaft pauschal und vorschnell abzuwerten. Aus der Kritik an der informationstheoretischen Musikanalyse sind alternative Analysemodelle entwickelt worden, deren Bewährungsprobe noch weitgehend aussteht, denen aber wenigstens eine Chance zur Bewährung eingeräumt werden sollte, was nur über den Versuch ihrer Anwendung und kritischen Weiterentwicklung geschehen kann (vgl. Abschnitt II. der Bibliographie).

Seit etwa zehn Jahren wird insbesondere auf dem Gebiet der Vergleichenden Musikwissenschaft (Ethnomusicology) in den USA, aber auch in Frankreich, der Sowjetunion und in Schweden, eine Darstellung musikalischer Strukturen mit dem methodischen Repertoire der Linguistik versucht. Dieser Ansatz konnte in einigen Veröffentlichungen auch schon bis zur expliziten Regelan-gabe für musikalische Strukturen vorangebracht werden (vgl. Abschnitt III. der Bibliographie). Die kritische Übernahme linguistischer Methoden und ihre Adaptation und Weiterentwicklung, die den strukturellen Besonderheiten der Musik und ihren Unterschieden zur Sprache gerecht werden, kennzeichnet die gegenwärtige Entwicklungsrichtung mathematischer Methoden im Bereich der Musikwissenschaft.

Das eingangs erwähnte Mißtrauen scheint in der Bundesrepublik besonders hartnäckig zu sein, denn außer einigen Präliminarien zu strukturellen Analogien zwischen Sprache und Musik (RAUHE, REINECKE, RIBKE, 1975) hat diese neue Entwicklung seit dem Nachlassen informationstheoretisch fundierter Analysen in der bundesrepublikanischen Musikwissenschaft noch keinen Widerhall gefunden. Die vorgelegte Bibliographie mag als Anstoß zu einer Entwicklung in diese Richtung verstanden werden.

Die nachfolgende Bibliographie enthält Zusammenfassungen der meisten Veröffentlichungen, die eine Anwendung quantitativer oder formal-analytischer Verfahren für Musik propagieren oder bereits mit diesen Methoden arbeiten. Einen Anspruch auf Vollständigkeit erhebt die Zusammenstellung dieser Literatur nicht.

Literatur

- de la Motte-Haber, H.: Musikpsychologie. Köln: Gerig, 1972
- de la Motte-Haber, H.: Psychologie und Musiktheorie.
Frankfurt: Diesterweg, 1976
- Rauhe, H., Reinecke, H.-P. & Ribke, W.: Hören und Verstehen.
München: Kösel, 1975
- Wiora, W.: Methodik in der Musikwissenschaft. In: M.Thiel (Hrsg.),
Methoden der Kunst- und Musikwissenschaft. Enzyklopädie
der geisteswissenschaftlichen Arbeitsmethoden, Bd. 6,
München: Oldenbourg, 1970, 93-139

BIBLIOGRAPHIE

I. INFORMATIONSTHEORETISCHE MUSIKANALYSE

1. Bauer, H.J.: Statistik, eine objektive Methode zur Analyse von Kunst? International Review of the Aesthetics and Sociology of Music, 1976, 7, 249-263.

Gegen die statistischen Analysemethoden von W. Fucks werden kritische Einwände erhoben, insbesondere wegen der Vernachlässigung der Zeitdimension in der Musik. Es wird als Beispiel eine statistische Analyse von G. Ligeti's "Lux aeterna" dargestellt, die nach Ansicht des Autors nichts erbracht hat, was nicht schon in der Musikgeschichte bekannt gewesen wäre, so daß Aufwand und Ergebnis der statistischen Musikanalyse in keinem akzeptablen Kosten-Nutzen-Verhältnis zueinander stehen.

2. Brincker, J.: Statistical analysis of music. An application of information theory. Svensk tidskrift för musikforskning, 1970, 52, 53-57

Der Aufsatz beschreibt die informationstheoretische Analyse und den Vergleich von 10 vierstimmigen Chorälen, die jeweils von drei verschiedenen Komponisten gesetzt wurden (Hassler, J.S. Bach, Weyse). Übergangswahrscheinlichkeiten 1., 2. und 3. Ordnung der Tonhöhen werden berechnet und zum Vergleich zwischen den drei Komponisten herangezogen.

3. Cohen, J.E.: Information theory and music. Behavioral Science, 1962, 7, 137-162

Im Rahmen eines Sammelreferats zur Anwendung informationstheoretischer Parameter als Indizes für musikalische Strukturen kritisiert der Autor die Übertragung der Informationstheorie auf die Menge "musikalischer Zeichen" unter zwei Gesichtspunkten: 1. Er zeigt, daß die mathematischen Voraussetzungen für einen stochastischen Prozeß, die Voraussetzung der Ergodizität, die Voraussetzung, daß die Quelle stationär ist und die Voraussetzung der Markoff-Konsistenz von musikalischen Zeichenstrukturen nicht erfüllt werden. 2. Er stellt eine Reihe von Argumenten dar, die

belegen, daß die psychologischen Voraussetzungen beim Hörer, die implizit bei der Übertragung der Informationstheorie auf die Musik gemacht werden, nicht erfüllt sind.

4. Coons, E. & Kraehenbuehl, D.: Information as a measure of structure in music. *Journal of Music Theory*, 1957, 1, 127-161

Die Autoren entwickeln ein neues Informationsmuster zur Beschreibung musikalischer Strukturen, das dem Prinzip gerecht wird, "that a significant work of art achieves 'unity' through 'variety'". "Unity" und variety" werden als Ähnlichkeits- bzw. Unähnlichkeitsbeziehungen einzelner Elemente in einem musikalischen Muster definiert, deren Häufigkeiten in die Informationsmuster eingehen. Ihr Index der "Artikuliertheit" ist ein Maß für die Ausgewogenheit von "unity" und "variety". Der Hierarchieindex ist ein Maß dafür, wie erfolgreich verschiedene Elemente so angeordnet werden, daß dabei der Eindruck der "unity" nicht verlorenght. (vgl. auch Kraehenbuehl, D. & Coons, E.: Information as a measure of the experience of music. *Journal of Aesthetics and Art Criticism*, 1959, 17, 510ff.)

5. Detlovs, V.K.: O statističeskom analize muzyki (Über die statistische Musikanalyse). *Latvijskij matematičeskij ežegodnik*. Riga: 1968, 3, 101-120

In allgemeiner Form werden hier die Methode und die Probleme der statistischen Analyse in der Musik behandelt, insbesondere geht es um die statistische Untersuchung von Melodie und Akkordik (Harmonik). Als konkretes Beispiel werden die Resultate einer vom Autor angestellten statistischen Harmonieanalyse der Choräle J.S. Bachs herangezogen. Es wird gezeigt, daß die Anwendung selbst einfacher quantitativer Charakteristika (beispielsweise Entropien erster und zweiter Ordnung) es im Rahmen einer solchen Untersuchung gestattet, nicht nur die allgemeinen Gesetzmäßigkeiten des Harmoniedenkens in Bach'schen Chorälen aufzuzeigen, sondern auch Unterschiede darin zu entdecken, wie Bach die Tongeschlechter Dur und Moll einsetzt, Tonalitäten, die sich hinsichtlich ihres emotionalen Charakters wesentlich unterscheiden.

6. Fucks, W.: Mathematische Musikanalyse und Randomfolgen. *Musik und Zufall*. *Gravesaner Blätter*, 1962, 6, 132-145

Statistische Methoden und Ergebnisse einer Häufigkeitsauszählung von Intervallen konsekutiver Töne werden dargestellt. Als geeigneter Kennwert zur Darstellung der Variation der Häufigkeitsverteilungen musikalischer Werke verschiedener Jahrhunderte stellt sich der Parameter Kurtosis heraus. Es zeigt sich, daß die Kurtosis für Werke des Vorbarock bis zu Werken der Spätromantik von 2,1 bis auf 12,1 ansteigt und bei zeitgenössischer Musik (Schoenberg, Witten, Berg, Henze, Stockhausen, u.a.) wieder auf den Wert 3,5 abfällt. Außerdem werden Ergebnisse von Korrelationsanalysen berichtet, die den Zufallsgehalt von Musikwerken quantitativ bestimmen. (Beispiele: Bach, Beethoven, Weber).

7. Fucks, W.: Gibt es mathematische Gesetze in Sprache und Musik? In: H. Frank (Hrsg.), *Kybernetik - Brücke zwischen den Wissenschaften*. Darmstadt, 1965, 220-234

Es wird durch Anwendung einer Reihe statistischer Kennwerte (Häufigkeiten und Übergangshäufigkeiten, Streuung, Kurtosis, Exzess und Korrelation) nachgewiesen, daß sich die Tonhöhenverteilungen in musikalischen Werken seit 1500 charakteristisch unterscheiden und daß es Gesetzmäßigkeiten in den Tonhöhenverteilungen in Werken einer musikalischen Epoche gibt. Damit wird anhand eines umfangreichen Datenmaterials demonstriert, daß eine quantitative Stilanalyse möglich ist und neue Einsichten in unterschiedliche Kompositionstechniken vermittelt. Dieselben Analysemethoden werden ebenfalls auf Sprache angewendet.

8. Fucks, W. & Lauter, J.: *Exaktwissenschaftliche Musikanalyse*. Köln/Opladen: Westdeutscher Verlag, 1965

Ziel der Untersuchungen ist die Erfassung der formalen Struktur von musikalischen Werken durch Darstellung quantitativer Parameter der Musik und deren statistischer Auswertung. Es wird versucht, Gesetzmäßigkeiten in der Variation dieser Parameter bei Vergleich verschiedener musikalischer Epochen zu entdecken. Z.B. wurde die Streuung und die Entropie von Tonhöhenverteilungen in

Werken aus der Zeit zwischen 1530 und 1960 untersucht und für beide Parameter ein Anstieg innerhalb dieses Zeitraums ermittelt. Es zeigt sich, daß die statistische Musikanalyse Charakteristika der Kompositionsstile einzelner Komponisten und ganzer musikalischer Epochen darzustellen vermag.

9. Heike, G.: Informationstheorie und musikalische Komposition. Melos, 1961, 28, 269-272

Von einer Darstellung informationstheoretischer Quantifizierungen musikalischer Strukturen ausgehend stellt der Autor die Anwendung einer "finite-state grammar" als Alternative zur Diskussion, die im Unterschied zu informationstheoretischen Parametern musikalische Strukturen unter Bezug auf die Positionen der Elemente innerhalb einer Zeichenkette beschreibt und übergeordnete Strukturen sichtbar werden läßt. Es wird die Möglichkeit angedeutet, solche Strukturbeschreibungen in Computerprogramme umzusetzen.

10. Hiller, L.A.: Informationstheorie und Musik. Darmstädter Beiträge zur Neuen Musik, Bd. 8, Mainz: Schott, 1964

Eingehende Darstellung der Methoden und Befunde eines Forschungsprogramms zur Anwendung der Informationstheorie auf die Musikanalyse, das vom Autor und Mitarbeitern am Studio für Experimentelle Musik an der Universität Illinois durchgeführt wurde. Es wurden vier Sonatenexpositionen (Mozart, Beethoven, Hindemith, Berg) bezüglich des Informationsgehalts auf der Grundlage von Tonhäufigkeiten und Tondauern miteinander verglichen. Außerdem wurden 16 Streichquartette (Mozart, Haydn, Beethoven) nach einer Vielzahl von Gesichtspunkten statistisch analysiert sowie eine umfassende statistische Analyse des 1. Satzes der Sinfonie op. 21 von Webern dargestellt. (Vgl. auch Hiller, L.A. & Bean, C.: Information theory analysis of four sonata expositions. Journal of Music Theory, 1966, 10, 96ff.)

11. Hiller, L. & Fuller, R.: Structure and information in Webern's symphony, op. 21. Journal of Music Theory, 1967, 11, 60-116

Der Artikel vergleicht die konventionellen strukturellen Analysen der Sinfonie op. 21 von Webern mit einer informationstheoretischen Analyse und zeigt, daß beide Analysemethoden sich gegenseitig ergänzen können. Die informationstheoretische Analyse besteht in der Berechnung von Redundanz und Informationsfluß basierend auf den Tonhöhenhäufigkeitsverteilungen und Intervallhäufigkeitsverteilungen. Es zeigt sich, daß die Tonhöhenverteilungen erster Ordnung einer Zufallsverteilung entsprechen.

12. Lewin, D.: Some applications of communication theory to the study of twelve-tone music. Journal of Music Theory, 1968, 12, 51-85

In diesem Artikel werden die Grundlagen einer Anwendung der Informationstheorie als analytische Methode für musikalische Werke dargestellt und mit Beispielen demonstriert. Anschließend wird die Anwendung auf Zwölfton-Musik diskutiert und eine informationstheoretische Analyse von Weberns Op. 25, Nr. 1 aufgrund von Übergangswahrscheinlichkeiten erster, zweiter und dritter Ordnung für Intervalle vorgenommen.

13. Mix, R.: Die Entropieabnahme bei Abhängigkeit zwischen mehreren simultanen Informationsquellen und bei Übergang zu Markoff-Ketten höherer Ordnung, untersucht an musikalischen Beispielen. Köln/Opladen: Westdeutscher Verlag, 1967, 39-80

Der Autor errechnete für die ersten Sätze dreier Streichquartette von Haydn jeweils für die 1. Violine Tonhöhenentropien, die auf ein zeitliches Raster bezogen wurden. Dabei zeigte sich, daß der Informationsfluß bei Betrachtung höherer N-gramme kleiner wurde. Wurde nur das unterschiedliche Auftreten der Töne pro Sekunde berücksichtigt, so ergab sich ein Wert von 60 bit/sec, die Auszählung von Hexagrammen lieferte dagegen einen Informationsfluß von 8 bit/sec.

14. Moles, A.: Informationstheorie der Musik. Nachrichtentechnische Fachberichte, 1956, 3, 47-55

Der Autor stellt die Grundlagen einer musikalischen Informationstheorie dar und führt dann die Unterscheidung zwischen semantischer und ästhetischer Information als zwei Aspekte des musikalischen Zeichenrepertoires ein: Die semantische Information ist die mit Worten wiedergebbare Anordnung, die transkribierbar und logisch ist, und die ästhetische Information ist die Anordnung, die den Spielraum zwischen der Komplexität und Freiheit der realen Botschaft und dem starren Schema der semantischen Botschaft ausnutzt, wie sie durch die Partitur vermittelt wird.

15. Pinkerton, R.C.: Information theory and melody. Scientific American, 1956, 194, 77-86

Auf der Grundlage von Auszählungen der Tonhöhenhäufigkeiten und Tonhöhenübergangshäufigkeiten erster Ordnung bei 39 amerikanischen Kinderliedern stellt der Autor eine Graphenstruktur zur Produktion künstlicher Kinderlieder dar (banal tune maker) und berechnet die Redundanz der analysierten und produzierten künstlichen Lieder. Die Relevanz einer möglichen Anwendung der Informationstheorie für die quantitative Analyse und die künstliche Produktion von Musik wird diskutiert.

16. Rojterštejn, M.I.: Graf i matrica kak instrumenty ladovogo analiza (Graph und Matrix als Instrumente der Tonstufenanalyse). In: Muzykal'noe iskusstvo i nauka. Moskva: Muzyka, 1973, 175-189

Bei dieser Arbeit handelt es sich um eine Untersuchung der Übergangswahrscheinlichkeiten von der Tonstufe i zur Stufe j in Themen J.S. Bachs (Fugen aus dem "Wohltemperierten Klavier"), D. Schostakowitsch's (Fugen op. 87) und W. Mozarts (Klaversonaten). Es wird gezeigt, daß die Übergangswahrscheinlichkeiten ein wichtiges Charakteristikum bei der statistischen Analyse von Werken "tonaler" Musik sein können. Der Autor betont, daß es möglich sei, eine Reihe wesentlicher Besonderheiten der Tonstufenstatistik in der Melodie auszumachen, indem man diese Wahrscheinlichkeiten mittels Matrizen darstellt.

17. Stahl, V.: Informationswissenschaft und Musikanalyse. Grundlagenstudien aus Kybernetik und Geisteswissenschaft, 1964, 5, 51-58

Es werden die Redundanzwerte für Tonhöhen- und Tondauernverteilungen für Werke aus der Gregorianik, aus dem Barock, aus der Klassik, der Romantik, der tonalen wie der atonalen Moderne und der "musica viva" berichtet. Für die Redundanzwerte der Tonhöhen und der Tondauernverteilungen zeigt sich eine kontinuierliche Abnahme zwischen 1500 und der Gegenwart.

18. Wagner, G.: Exaktwissenschaftliche Musikanalyse und Informationsästhetik. International Review of the Aesthetics and Sociology of Music, 1976, 7, 63-76

Diese Arbeit setzt sich kritisch mit den Methoden und Ergebnissen der exaktwissenschaftlichen Musikanalyse nach W. Fucks auseinander.

19. Winkel, F.: Die informationstheoretische Analyse musikalischer Strukturen. Musikforschung, 1964, 17, 1-14

Ausführliche Darstellung der Anwendungsgesichtspunkte der Informationstheorie für eine quantitative Musikanalyse mit Darstellungen von Tonhöhenverteilungen (Schönberg "Das Buch der hängenden Gärten" und Viktor Holländer "Schön war's doch").

20. Youngblood, J.E.: Style as information. Journal of Music Theory, 1958, 2, 24-35

Ziel der Untersuchung ist eine Beurteilung der Angemessenheit einer informationstheoretischen Beschreibung musikalischer Stile. Der Autor analysierte 8 Lieder aus Schuberts "Die Schöne Müllerin", 8 Arien von Mendelssohn und 6 Lieder aus Schumanns "Frauen - Liebe und Leben" hinsichtlich der Häufigkeit des Auftretens der 12 Töne der chromatischen Tonleiter. Es wurden Übergangswahrscheinlichkeiten 1. Ordnung berechnet und auf dieser Basis die Entropie, relative Entropie und die Redundanz berechnet. Es zeigte sich, daß für die Mendelssohn-Arien die höchste Redundanz, für die Schumann-Lieder die zweithöchste und für die Schubert-Lieder die geringste Redundanz festzustellen war.

II. ANDERE QUANTITATIVE UND FORMALE ANALYSEMETHODEN

21. Adams, C.R.: Melodic contour typology. Ethnomusicology, 1976, 20, 179-215

Es wird eine Typologie von Melodiekonturelementen und Melodiekonturen höherer Ordnung entwickelt und 30 Parameter durch quantitative Beschreibung von Melodiekonturen definiert. Diese werden zum Vergleich zweier Gruppen von Folk Songs herangezogen, die Ähnlichkeiten und Unterschiede bezüglich der Melodiekonturen dieser beiden Gruppen von Folk Songs quantitativ zu erfassen erlauben.

22. Boroda, M.G.: O melodičeskoj elementarnoj edinice. (Über eine elementare melodische Einheit). In: MAA-FAT-75. Materialy Pervogo Vsesojuznogo seminaru po mašynnym aspektam algoritmičeskogo formalizovannogo analiza muzykal'nych tekstov. Erevan: Izd. Armjanskoj SSR 1977, 112-120

In dieser Arbeit wird die zum erstenmal vom Autor isolierte, eindeutig bestimmte, elementare metrisch-rhythmische Melodieeinheit "Formales Motiv" (F-Motiv) beschrieben, deren wesentliche Eigenschaften hier analysiert werden. Gezeigt wird, daß sich jede melodische Linie mit Taktstruktur von Anfang bis Ende lückenlos in F-Motive gliedern läßt; diese Segmentierung ist wesentlich von der metrisch-rhythmischen Struktur der melodischen Linie abhängig, d.h. von den rhythmischen und metrischen Wechselbeziehungen ihrer Töne. Der Autor unterstreicht, daß das F-Motiv eine natürliche elementare Einheit in den Melodien eines breiten Kreises musikalischer Stile und ein wirkungsvoller Parameter bei der statistischen Melodieanalyse ist. Es wird gezeigt, daß das F-Motiv hinsichtlich gewisser quantitativer Charakteristika (u.a. der mittleren Länge) dem Wort ähnlich ist. Aus Werken unterschiedlicher Stilrichtungen sind Beispiele für die Gliederung melodischer Fragmente in F-Motive angeführt.

23. Boroda, M.G., Nadarejšvili, F.S., Orlov, J.K., Čitašvili, R.J.: O karaktere raspredelenija informacionnyh edinic maloju častoty v chudožestvennyh tekstach. (Über den Verteilungscharakter informatorischer Einheiten mit geringer Häufigkeit in künstlerischen Texten). In: Semiotika i informatika 9, 1977, 23-34

Anhand literarischer sowie musikalischer Texte unterschiedlicher Stilrichtungen wird in der Arbeit die Verteilung selten vorkommender Elemente im künstlerischen Text analysiert. In den literarischen Texten ist als Element das Wort gewählt, in den musikalischen die elementare Melodieeinheit "Formales Motiv" (F-Motiv). Vgl. die unter 22. genannte Arbeit: M.G. Boroda, Über eine elementare melodische Einheit.

Es wird gezeigt, daß selten vorkommende Elemente eine deutliche Tendenz zur Häufung, d.h. zur Konzentration auf kleine Textabschnitte zeigen. Außerdem wird ein quantitatives Maß für den Grad dieser Häufung vorgeschlagen. Die Autoren unterstreichen, daß die gewonnenen Resultate auf eine klare statistische Heterogenität zwischen literarischem und musikalischem Text hindeuten. In der Anlage zur Arbeit werden neben der formalen Bestimmung des F-Motivs auch die Daten angeführt, die sich beim Vergleich der Verteilung von F-Motivlängen im musikalischen Text mit der Verteilung von Wortlängen im literarischen Text ergeben haben.

24. Boroda, M.G.: Ob opredelenii informacionnoj melodičeskoj edinicy tipa frazy v muzyke (Über die Bestimmung einer informatorischen Melodieeinheit vom Typ der Phrase in der Musik). In: Soobščeniya AN GSSR 89, 1978, 57-60

Ausgehend von den in der Musiktheorie bekannten Prinzipien der rhythmischen Kohäsion von Tönen wird in der Arbeit zum erstenmal eine eindeutig bestimmte große Melodieeinheit isoliert. Es wird gezeigt, daß diese Einheit, "Rhythmische Phrase" (R-Phrase) genannt, in einer Vielzahl musikalischer Stilrichtungen anzutreffen ist. Veranschaulicht wird die Effizienz der Melodieanalyse unter Verwendung des Konzepts der R-Phrase mittels quantitativer Methoden, insbesondere bei der Untersuchung von Unterschieden zwi-

schen instrumentaler und vokaler Musik.

25. Boroda, M.G., Orlov, Ju.K.: O sravnenii skorostej vospri-
jatija informacijonnych edinic v muzyke i reči (Ver-
gleich des Wahrnehmungstempos bei Informationseinheiten
in Musik und Sprache). In: Materialien der Konferenz
"Das Funktionieren von Kommunikationssystemen". Erevan:
Akademie der Wissenschaften der Armenischen SSR, 1978,
58-62

Die Untersuchungsfrage ist: Mit welchem Tempo werden kleine Ele-
mente in literarischen und musikalischen Werken bei der Aufführung
aufgefaßt? Als Material wurden Schallplatten von Werkaufführungen
durch die Autoren oder durch bekannte Darsteller verwendet. Das
sprachliche Material bestand aus Poesie und Prosa verschiedener
Autoren und enthielt auch Ausschnitte aus Schauspielen (alle
Stücke waren in Russisch verfaßt). Das musikalische Material be-
stand aus Kompositionen verschiedener Stilformen und Gattungen;
es enthielt sowohl instrumentale (Sonate, Symphonie, Präludium
und Fuge, Rondo) als auch Vokalkompositionen (Romanzen und Opern-
arien, in Russisch). Als Elementareinheit (kleines Element) galt
in den literarischen Texten das Wort, in den musikalischen Tex-
ten die elementare melodische Einheit des "formalen Motivs" (F-
Motiv; M.G. Boroda: K voprosu o metroritmičeski elementarnoj
edinice v muzyke (Zur Frage einer metrorhythmisch elementaren
Einheit in der Musik). Soobščeniya AN GSSR, 1973, 71, Nr. 3).
Es zeigte sich, daß die Auffassungsgeschwindigkeit von Wörtern
bei der Aufführung eines literarischen Werks im Durchschnitt der
Auffassungsgeschwindigkeit von F-Motiven bei der Aufführung musi-
kalischer Werke sehr ähnlich ist. Ferner zeigte sich, daß in Vo-
kalkompositionen die Gesamtzahl von F-Motiven der Melodie im all-
gemeinen ganz nahe mit der Wortzahl des zur Melodie gehörenden
Textes übereinstimmt. Bezüglich dieser Befunde werden einige Er-
klärungshypothesen vorgeschlagen.

26. Crane, F., Fiehler, J.: Numerical methods of comparing
musical styles. In: H.B. Lincoln (Ed.), The Computer and
Music. Ithaca: Cornell University Press, 1970, 209-222

Der Artikel bespricht den Gebrauch von Computern bei Vergleichen
von musikalischen Stilen. Die verschiedenen Vergleichstechniken
sind nach ihrer logischen Aufeinanderfolge in der Prozedur geord-
net:

1. Rohdaten und ihre numerische Darstellung: Die Daten für ei-
nen Vergleich werden in einer Matrix zusammengestellt. Sie sollen
möglichst elementare Charakteristika der einzelnen Musikwerke
sein. Danach werden sie miteinander korreliert.

2. Messung der Affinität zwischen 2 Werken: Nach SOKAL und
SNEATH kann man die Methoden, mit denen Affinität gemessen wird,
in 3 Klassen einteilen, - Assoziation, Korrelation und Distanz.
Zwischen den musikalischen Charakteristika zweier Werke werden
Koeffizienten berechnet, die dann aussagen, in welchem Grad zwei
Musikwerke assoziiert, miteinander korreliert oder voneinander
entfernt sind.

3. Methoden zum 'Clustern' von Stilen: Das Computerprogramm
clustert verschiedene Musikwerke nach vorher festgelegten Kri-
terien. Die verschiedenen Kriterien bestimmen die verschiedenen
Methoden.

Die oben dargestellten Methoden werden an einem Beispiel ver-
deutlicht und weitere mögliche Methoden der multivariaten Sta-
tistik mit ihren Anwendungsmöglichkeiten aufgezählt.

27. Erickson, R.: General purpose system for computer-aided music-
al studies. Journal of Music Theory, 1969, 276-294

Erickson stellt in seinem Artikel ausführlich ein System dar,
(general-purpose-compiler), das in der Lage ist, da es auf ein-
zelnen syntaktischen Eigenschaften beruht, jede musikalische Kom-
position zu analysieren.

28. Heckmann, H. (Hrsg.): Elektronische Datenverarbeitung in der
Musikwissenschaft. Regensburg: Bosse, 1957

Der Sammelband enthält eine Reihe von Aufsätzen zur computerun-

terstützten quantitativen Musikanalyse sowie eine Bibliographie zu diesem Forschungsgebiet. Die abgehandelten Themen reichen von allgemeinen Problemen der Datenverarbeitung im Bereich der Musikwissenschaft über die Darstellung von Programmsystemen zur quantitativen Musikanalyse bis zu Anwendungsbeispielen bei musikalischen Werken.

29. Iglickij, M.A.: Rodstvo tonal'nostej i zadača ob otyskanii moduljacionnych planov (Die Verwandtschaft von Tonalitäten und die Aufgabe des Auffindens von Modulationsplänen). In: Muzykal'noe iskusstvo i nauka, II. Moskva: Muzyka, 1973, 192-205

Eine für die Musikanalyse "mit exakten Methoden" wichtige Aufgabe wird einer strikten Lösung zugeführt: Es wird ein Verfahren vorgeschlagen und ein Algorithmus explizit beschrieben, mit dem für zwei beliebige Tonalitäten des Dur-Moll-Systems alle Wege minimaler Länge für die schrittweise Modulation der einen in die andere aufgefunden werden können. Der Algorithmus wurde vom Autor auf der Grundlage eines von ihm entworfenen Systems mathematisch strikter Begriffe konstruiert, mit denen die Tonalitätsverwandtschaft im klassischen Dur-Moll-System zusammenfassend und explizit beschrieben werden kann. Die Begriffe entsprechen den in der Musiktheorie akzeptierten Vorstellungen. Die Arbeit ist von Bedeutung für die Untersuchung der Tonart und der Tonalität mittels exakter Methoden. Ihre Ergebnisse können effektiv zur Analyse von Kompositionsstilen eingesetzt werden (Untersuchung der Modulationsstatistik und der Modulationspläne usw.).

30. Karlina, T.P., Detlovs, V.K.: Statističeskij analiz sekvencij v melodii (Statistische Analyse von Sequenzen in der Melodie). In: Latvijskij matematičeskij ežegodnik. Riga 1968, 4, 165-173

Die Autoren haben anhand von Melodien J.S. Bachs, J. Haydns und P. Tschaikowskys eine statistische Analyse von Sequenzen, parallelen Verschiebungen kleiner abgeschlossener Strukturen in Melodie oder Harmonie, angestellt. In Abhängigkeit von Intervall und Tonstufe wurden die Sequenzen klassifiziert, und es wurde ge-

zeigt, daß die Anwendung statistischer Kennwerte (Mittelwert und Streuung) es gestattet, nicht nur die kennzeichnenden Züge sowie die individuellen Besonderheiten hinsichtlich der Nutzung von Sequenzen bei jedem dieser drei Komponisten aufzuzeigen, sondern daneben auch einige allgemeine Evolutionstendenzen für diese Verwendung im Entwicklungsprozeß der Musik nachzuvollziehen.

31. Kolinski, M.: The general direction of melodic movement. Ethnomusicology, 1965, 9, 240-264

Der Autor beschreibt eine quantitative Methode zur Analyse der vorherrschenden Richtung melodischer Bewegungen. Die Analyse-methode wird zum interkulturellen Vergleich herangezogen, indem Häufigkeitsverteilungen der beschriebenen deskriptiven Parameter miteinander verglichen werden.

32. Rohrer, H.G.: Musikalische Stilanalyse auf der Grundlage eines Modells für Lernprozesse. Berlin/München: Siemens AG, 1970

Auf der Grundlage eines Modells des Verstehens von Musik, d.h. der kognitiven Analyseprozesse des Musikhörers, entwirft der Autor ein mathematisches Modell zur quantitativen Analyse musikalischer Werke, das im Unterschied zu herkömmlichen Verfahren (z.B. Informationstheorie) die Möglichkeiten und Grenzen des Empfängers mitberücksichtigt. In dieses Modell gehen insbesondere Annahmen über psychische Prozesse bei der Erwartungsbildung und der musikalischen Mustererkennung ein. Das auf der Basis dieses Modells geschriebene Computerprogramm zur Musikanalyse gestattet die quantitative Kennzeichnung einzelner musikalischer Werke sowie Ähnlichkeitsvergleiche zwischen verschiedenen Werken, die stilistische Kategorisierungen ermöglichen.

33. Rothgeb, J.: Some Uses of Mathematical Concepts in Theories of Music. Journal of Music Theory, 1966, 10, 200-215

Der Aufsatz gibt einen knappen Überblick über die Anwendung mathematischer Methoden bei der Beschreibung von 12-Ton-Systemen, nicht-serieller, atonaler Musik und bei der Entwicklung axiomatischer

tischer Systeme im Bereich der musiktheoretischen Forschung. Das primäre Ziel mathematischer Modelle für die 12-Ton-Musik ist die Charakterisierung von Transformationen und Invarianzen bei Anwendung mathematischer Operationen. Ähnlichkeitsrelationen bei nicht-serieller, atonaler Musik werden durch Tonhöhenmengen und Mengenoperationen formalisiert. Anschließend wird auf Probleme bei der Anwendung der mathematischen Logik bei der Entwicklung generativer Systeme eingegangen. Weitere Arbeiten zu diesem Themenbereich kann man der Bibliographie dieses Aufsatzes entnehmen.

34. Selleck, J., Bakeman, R.: Procedures for the analysis of form. Two computer applications. Journal of Music Theory, 1965, 9, 281-295

Die Autoren wollen in ihrem Artikel zeigen, wie hilfreich Datenverarbeitungsprozeduren bei der Bearbeitung komplizierter musikalischer Analysen sind. Als Beispiel nehmen sie Gregorianische Gesänge, die mit Hilfe verschiedener Programme nach 3 formalen Kriterien analysiert werden. Die verschiedenen Möglichkeiten und die Ergebnisse der Analysen (und deren kritische Betrachtung) sind ausführlich dargestellt.

35. Simon, H.A., Sumner, K.: Pattern in music. In: B. Kleinmuntz (Ed.), Formal Representation of Human Judgement. New York: Wiley, 1968, 219-251

Es wird in diesem Artikel ein Schema aufgezeigt, nach dem Musik durch Muster beschrieben wird. Außerdem wird eine Sprache (formal/mathematisch) dargestellt, durch die Muster beschrieben werden können, und an Beispielen ihre Anwendung gezeigt. Um diese Sprache zu testen, werden 2 Computerprogramme vorgestellt, von denen das eine Musterbeschreibungen in Musik, das andere die Noten in Musterbeschreibungen übersetzen soll.

36. Smoliar, S.W.: Music programs: An approach to music theory through computational linguistics. Journal of Music Theory, 1976, 20, 105-131

Der Autor versucht, die Konzepte "syntaktische", "phonologische"

und "semantische" Komponente einer Sprache auf Programmiersprachen zu übertragen, die der Analyse und Produktion von Musik mit Hilfe eines Computers dienen sollen. Auf dieser Basis wird ein Modell interagierender Kontrollprozesse entworfen, die ein Computer bei der Analyse und Produktion von Musik auszuführen hat. Dieses Modell kann als Analogie zu einem Modell der beim menschlichen Hören ablaufenden kognitiven Analyseprozesse aufgefaßt werden.

III. ÜBERTRAGUNG LINGUISTISCHER METHODEN AUF DIE STRUKTURELLE ANALYSE VON MUSIK

37. Aranovskij, M.G.: Myšlenije, jazyk, semantika (Denken, Sprache, Semantik). In: Problemy muzykal'nogo myšlenija. Moskva: Muzyka 1974, 90-128

In der Arbeit wird der Begriff der "Musiksprache" von semiotischen Positionen her einer eingehenden Betrachtung unterzogen. Angeführt werden eine Reihe wesentlicher Besonderheiten, die die "musikalische Sprache" von der verbalen unterscheiden. Angeschnitten wird das Problem der Isolierung von semantischen Einheiten im musikalischen Werk. Daneben wird der Versuch unternommen, angefangen bei den elementaren Segmenten bis hin zu den großen Einheiten vom Typ der musikalischen Phrase, eine Bestimmung des Systems solcher Einheiten zu liefern. Beginnend mit dem allgemeinen Plan und den allgemeinen Proportionen der Konstruktion bis hin zu seiner konkreten tonalen Verwirklichung wird ein vom Autor vorgeschlagenes Erzeugungsschema (-modell) des musikalischen Werkes beschrieben.

38. Arom, S.: Essai d'une notation des monodies à des fins d'analyse. Revue de musicologie, 1970, 55, 172-216

Der Autor wendet die von RUWET vorgeschlagene strukturalistische Analysemethode auf Musik nicht-westlicher Kulturen an. Es werden Distributionen und Oppositionen einer Vielzahl musikalischer Merkmale dargestellt.

39. Boillès, C.: Tepehua-thought-song: A case of semantic signalling. *Ethnomusicology*, 1967, 11, 267-292

Es wird eine Transformationsgrammatik für Lieder der Tepehuas dargestellt. Diese Lieder sind mit einem gesprochenen linguistischen Kode vergleichbar, weil bestimmten Tonfolgen spezifische Bedeutungen gegeben werden. Die Transformationssyntax macht die Zusammenhänge zwischen rhythmischer und melodischer Struktur und der Bedeutung der Tonfolge deutlich.

40. Chenoweth, V., Bee, D.: Comparative-generative models of a New Guinea melodic structure. *American Anthropologist*, 1971, 73, 773-782

Die Autoren stellen mit Hilfe von Flußdiagrammen, Formeln und geometrischen Modellen melodische Strukturen dar, die die Basis für generative Regeln einer Transformationssyntax liefern. Das Ziel der Analyse ist es, alle syntaktisch korrekten Melodien, die die Syntax produzieren kann, vorherzusagen. (vgl. auch Chenoweth, V.: *Melodic perception and analysis*. Papua, New Guinea: Summer Institute of Linguistics, 1972).

41. Cooper, R.: Abstract structure and the Indian Rāga system. *Ethnomusicology*, 1977, 21, 1-32

Der Autor beschreibt die Methode der Aufstellung einer generativen Syntax für Melodiestrukturen und stellt eine generative Transformationssyntax für die Melodien indischer Rāgas dar (vgl. auch Cooper, R.: *Propositions pour un modèle transformationnel de description musicale*. *Musique en jeu*, 1973, 10, 70-88).

42. Feld, S.: Linguistic models in ethnomusicology. *Ethnomusicology*, 1974, 18, 197-217

Der Autor kritisiert den der Linguistik analogen Ansatz der musikalischen Analyse im Bereich der Vergleichenden Musikwissenschaft. Einer der Hauptkritikpunkte ist der, daß die Vergleichende Musikwissenschaft nicht unkritisch die Syntaxmodelle der Linguistik übernehmen kann, weil einmal die Analogie von Sprache und Musik, die die grundlegende Voraussetzung für diesen Ansatz bildet, bis-

her nur unzureichend belegt worden sei und weil eine Vielzahl anderer Voraussetzungen, die die linguistischen Modelle machen (Kompetenzmodell, Tiefenstruktur als Bedeutungsträger, Unabhängigkeit von kulturellen Bedingungen, Problem des idealen Hörers, Bewertungsprinzipien für die Akzeptierbarkeit eines Syntaxmodells), nicht reflektiert werden. Der Autor betont, daß es für die Vergleichende Musikwissenschaft von primärer Bedeutung ist, die kulturellen Bedingungen zu erklären, die Musik syntaktisch akzeptierbar werden läßt.

43. Gasparov, B.M.: O nekotorych principach strukturnogo analiza muzyki (Über einige Prinzipien der strukturellen Musikanalyse). In: *Problemy muzykal'nogo myšlenija*. Moskva: Muzyka 1974, 129-152

In dieser Arbeit wird ein Zugang zur Musikanalyse dargelegt, der von den Positionen der strukturellen Linguistik ausgeht. Beschrieben wird ein vom Autor vorgeschlagenes Konstruktionsschema für ein Modell der Musiksprache. Als konkrete Anwendung strukturell-linguistischer Methoden bei der Untersuchung von Musik wird die Distributionsanalyse der Harmoniesprache Beethovens (frühe Schaffensperiode) dargestellt.

44. Lévi-Strauss, C.: "Bolero" de Maurice Ravel. *L'Homme*, 1971, 11, 5-14

Der Autor versucht eine Analyse des Boleros von Ravel mit den Methoden der strukturalistischen Linguistik.

45. Nattiez, J.J.: Is a descriptive semiotics of music possible? *Language Sciences*, 1972, 23, 1-7

Von einem Vergleich zwischen Sprache und Musik, insbesondere von Phonem und Note, ausgehend belegt der Autor, daß ein funktionalistischer Zugang zur Semiotik der Musik nicht möglich ist, weil Noten im Unterschied zu Phonemen nicht zu einem semantischen System gehören, das erst eine Unterscheidung phonologischer Einheiten durch Bedeutungsunterschiede auf der Ebene der "first articulation" erlaubt. Stattdessen schlägt der Autor eine strukturalistische Vorgehensweise vor, die in einem algorithmischen An-

satz bestehen müßte, der es gestattet, musikalische Einheiten zu segmentieren, zu klassifizieren und Transformationen von Einheiten zu identifizieren.

46. Nattiez, J.J.: Linguistics: A new approach for musical analysis? International Review of the Aesthetics and Sociology of Music, 1973, 4, 51 - 69

Der Autor stellt die Methodologie und die bei ihrer Anwendung auftretenden Probleme einer die Methoden der Linguistik übernehmenden Musikwissenschaft in ihren Grundzügen dar, wobei er der strukturalistischen Vorgehensweise den Vorzug einräumt.

1. Die Übertragung linguistischer Methoden auf die Musikanalyse läßt zu einer präzisen Definition der Analysekonzepte und -ziele ein. 2. Sie macht es notwendig, wiederholbare analytische Methoden zu entwickeln, die es gestatten, musikalische Codes zu entdecken. 3. Sie läßt es notwendig erscheinen, eine bislang nicht entwickelte wissenschaftliche Sprache zu verwenden, um strukturelle Beschreibungen zu ermöglichen. 4. Sie weist uns den Weg, durch den die Analyseergebnisse validiert werden können.

47. Nettle, B.: Some linguistic approaches to musical analysis. Journal of the International Folk Music Council, 1958, 10, 37 - 41

Der Aufsatz diskutiert Methoden der deskriptiven Linguistik unter dem Gesichtspunkt ihrer Übertragbarkeit auf die Musikanalyse. Kernpunkt der Diskussion ist die Frage der Übertragbarkeit der Methode der Distributionsanalyse auf musikalische Elemente. Dabei ergibt sich das Problem, daß bei größer werdendem Corpus die Distribution musikalischer Elemente nicht mehr in Form von Regeln dargestellt werden kann, weil die Zahl der Ausnahmen immer weiter ansteigt. Stattdessen schlägt der Autor eine statistische Analyse und Interpretation der Distributionen vor. Als Beispiel einer Distributionsanalyse teilt der Autor die Distribution von Intervallen auf betonten vs. unbetonten rhythmischen Positionen aus dem Anfang von Ernst Kreneks Op. 83, Nr. 1 mit.

48. Rojterštejn, M.I.: "Muzykal'nye" struktury v poezii Bloka ("Musikalische" Strukturen in der Poesie Bloks). In: O Muzyke. Problemy analiza. Moskva: Sovetskij kompozitor, 1974, 323 - 343

Kleine Strukturen mit geschlossener Bedeutung in Musik und Poesie werden auf Ähnlichkeiten ihrer inneren Organisation untersucht. In der Poesie handelt es sich um den Vierzeiler, in der Musik um die Periode. Untersucht werden Gedichte von A. Blok und Ausschnitte aus verschiedenen Musikwerken. Der Autor zeigt, wie verbreitet in Bloks Poesie (in der Struktur der Vierzeiler) metrische Beziehungen sind, die für den inneren Aufbau der achttaktigen Periode innerhalb verschiedener musikalischer Stilrichtungen charakteristisch sind, Beziehungen, die in der Musikwissenschaft gut bekannt sind und früher für "rein musikalisch" gehalten wurden. Außer diesem Hauptresultat werden in der Arbeit interessante Definitionen der Zäsurtiefe (Artikulationsstärke) im Vierzeiler gegeben, die sich auf die grammatische Struktur stützen. Die Arbeit ist zweifellos sowohl für die Musikwissenschaft als auch für die Sprach- und Literaturwissenschaft interessant und perspektivenreich.

49. Ruwet, N.: Language, Musique, Poésie. Paris 1972

Der Band enthält mehrere Aufsätze des Autors zum Thema der musikalischen Strukturanalyse. Besondere Beachtung verdienen im Zusammenhang mit dieser Bibliographie die beiden letzten Aufsätze, die den Grundlagen einer syntaktischen Musikanalyse gewidmet sind. Hier werden erste Schritte einer von der Linguistik beeinflussten bzw. zu ihr hinführenden Analysemethode dargestellt, die in einer Segmentierung auf verschiedenen Ebenen beruht, wobei nach Wiederholungen zumindest analoger Art gesucht wird, deren Distribution in Beziehung zu anderen Segmenten dargestellt werden kann. Der Autor wendet den Ansatz insbesondere auf Musik des 14. Jahrhunderts an, gesteht aber auch ein, daß die Komplexität romantischer Musik auf diese Art nicht adäquat erfaßt werden kann.

50. Sapir, J.D.: Diola-Fogny funeral songs and the naive critic. African Language Review, 1969, 8, 167 - 191

Der Autor stellt eine Syntax von Beerdigungsgesängen der Diola-Fogny dar, in der Phrasenstrukturregeln allen Gesängen gemeinsame Strukturmerkmale und Transformationsregeln mögliche Variationen erzeugen.

51. Sbornik "Problemy taksonomii estonskich runičeskich melodij" (Probleme der Taxonomie estnischer Runenmelodien). Tallin: Akademie der Wissenschaften der Estnischen SSR, 1977

Der Sammelband enthält Artikel zur Analyse estnischer Runenmelodien mit Hilfe der Methoden der strukturalen Linguistik. Die Beziehung zwischen Text und Melodie der Lieder wird im Detail untersucht. Auch die Anwendung formaler Grammatiken (insbesondere derjenigen von Chomsky und Lindenmaier) zur Modellierung der Melodien von Runenliedern wird erörtert. Es wird untersucht, ob sich der grammatische Zugang für Melodien mit geringen Unterschieden eignet.

52. Sbornik "Točnye metody i muzykal'noe iskusstvo" (Exakte Methoden und musikalische Kunst). Materialien zu einem Symposium. Rostov: Rostover Universitätsverlag, 1972

Der Sammelband enthält die auf dem 1972 in Rostov / Don durchgeführten Symposium "Exakte Methoden und musikalische Kunst" vorgelegten Materialien - Vorträge, Mitteilungen, Thesen. Sie umfassen einen weiten Kreis von Problemen, die mit der Anwendung formaler, insbesondere mathematischer Methoden bei der Untersuchung von Musik verbunden sind: methodologische Probleme (sie betreffen die Logik des musikalischen Denkens und die Möglichkeiten, den Kompositionsprozeß mit formalen Methoden zu untersuchen), experimentelle Untersuchungen statistischer Merkmale von Kompositionsstilen, Arbeiten zu Modellierung verschiedener Kompositionen auf EDV-Anlagen, Arbeiten zur musikalischen Akustik und semiotische Untersuchungen eines musikalischen Werks.

53. Springer, G.P.: Language and music: Parallels and divergencies. In: M. Halle, H.G. Lunt, H. McLean & C.H. Van Schooneveld (Eds.), For Roman Jakobson. Den Haag: Mouton, 1956, 504 - 513

Indem der Autor Ähnlichkeiten und Unterschiede zwischen Sprache und Musik aufzeigt, stellt er die Bedeutung einer Übertragung linguistischer Forschungsmethoden auf den Bereich der Musikanalyse dar, insbesondere informationstheoretischer und statistischer Analysen im Kontext interkultureller Vergleiche von Volksmusik.

54. Stoffer, T.H.: Aspekte einer generativen Syntax zur Beschreibung musikalischer Strukturen für eine kognitive Musikpsychologie. Berichte des Psychologischen Instituts der Ruhr-Universität Bochum, Arbeitseinheit Kognitionspsychologie, Nr. 11, 1979

Der Bericht stellt in Form expliziter und formaler Regeln eine generative Transformationssyntax für einen Corpus deutscher Volks- und Kinderlieder vor. Probleme der deskriptiven und kognitiven Adäquatheit struktureller Analysemethoden für Musik werden diskutiert. Der heuristische Nutzen einer musikalisch-syntaktischen Strukturanalyse wird durch die Beschreibung eines laufenden Forschungsprojekts darzulegen versucht, in dem eine experimentell-psychologische Analyse kognitiver Prozesse beim Hören, Lernen und Behalten von Melodien vorgenommen wird.

55. Sundberg, J., & Lindblom, B.: Generative theories in language and music descriptions. Cognition, 1976, 4, 99 - 122

Der Stil schwedischer Kinderlieder wird als ein generatives Regelsystem beschrieben. Ausgehend von einer Distributionsanalyse von Tonhöhen, Intervallen, Harmonien, Pausen u. dgl. wird ein System generativer Regeln teilweise explizit formuliert, das zur Produktion neuer "Kinderlieder" verwendet wird. Das Regelsystem lehnt sich eng an CHOMSKY & HALLE (1968) an. Die Autoren nehmen an, daß die von ihnen aufgewiesenen Parallelen zwischen Sprache und Musik typische kognitive Kapazitäten des Menschen reflektieren. (Vgl. auch Lindblom, B. & Sundberg, J.: Towards a generative

theory of melody. Svensk Tidskrift för Musikforskning, 1970, 52, 71 - 88)

56. Winograd, T.: Linguistics and computer analysis of tonal harmony. Journal of Music Theory, 1968, 12, 2 - 49

Der Artikel beschreibt eine Syntax für die Harmoniefolgen tonaler Musik, die die Grundlage für ein Programm zu Identifikation und Beschreibung harmonischer Strukturen darstellt. Ein interessanter Aspekt dieser Syntax besteht in der Auflösung von harmonischen Mehrdeutigkeiten durch "semantische" Analysen (harmonische Funktionsanalyse im Kontext der Komposition).

CURRENT BIBLIOGRAPHY

GENERAL

1. ALTMANN, G.: Towards a theory of language. Altmann, G. (Ed.): Glottometrika 1, Bochum: Brockmeyer, 1978, 1-25.
2. ALTMANN, G.: Prolegomena to Menzerath's law. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer, 1980, 1-10.
3. BAILEY, R.W.: The future of computational linguistics. ALCC Bulletin 7, 1979, 4-11.
4. BRAINERD, B.: On the Markov nature of text. Linguistics 176, 1976, 5-30.
5. ERMOLENKO, G.V., MOGILEVSKIJ, R.I.: Učebnoe posobie po lingvističeskoj statistike dlja studentov-filologov [A Textbook of Linguostatistics for Students of Philology]. Samarkand: Samarkandskij gosudarstvennyj universitet, 1979, 101 pp.
6. GROTHJAHN, R.: The theory of runs as an instrument for research in quantitative linguistics. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer, 1980, 11-43.
7. HOFFMANN, L., PIOTROWSKI, R.G.: Beiträge zur Sprachstatistik. Leipzig: VEB Verlag Enzyklopädie, 1979, 214 pp.
8. ITKONEN, E.: Qualitative vs. quantitative analysis in linguistics. Perry, Th.A.: (Ed.): Evidence and Argumentation in Linguistics. Berlin - New York: de Gruyter, 1980, 334-366.
9. IVANOV, V.S.: Količestvennye metody v sovremennom vengerskom jazykoznanii (1973-1975) [Quantitative methods in modern Hungarian linguistics (1973-1975)]. Voprosy jazykoznanija 1979/6, 121-130.
10. JAKUBAJTIS, T.A., BAXAREV, A.T.: Dispersionnyj analiz v lingvostatistike [Dispersion analysis in linguostatistics]. Izvestija Akademii nauk Latvvijskoj SSR 1979/1, 114-123.

25. SCHWARTZ, L.J.: Syntactic markedness and frequency of occurrence. Perry, Th.A. (Ed.): Evidence and Argumentation in Linguistics. Berlin - New York: de Gruyter, 1980, 315-333.

SEMANTICS

26. MARX, W.: Die Dominanz des Substantivs als Träger der assoziativen Bedeutung. Zeitschrift für experimentelle und angewandte Psychologie 26, 1979, 596-602.
27. McKINNON, A.: Relations between the most common Danish words. Sprache und Datenverarbeitung 2, 1978, 158-170.

LEXICOLOGY

28. ANDREEV, N.D.: Verteilungswörterbücher und Grundwortschatz. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer, 1980, 82-88.
29. GAVRILOVA, A.I.: Strukturno-verojatnostnye xarakteristiki francuzskoj sportivnoj leksiki: po dannym raspredelitel'nogo slovarja [The structural probabilistic characteristics of French sports-lexic: from the data of the distributive dictionary]. Lingvističeskie issledovanija 1978. Problemy leksikologii, leksikografii i prikladnoj lingvistiki. Moskva 1978, 30-36.
30. GAVRILOVA, A.I.: Rezul'taty statističeskoj obrabotki leksiki pervogo urovnja raspredelitel'nogo slovarja /RS/ russkogo jazyka i sopostavlenie dannyx RS russkogo i francuzskogo jazykov [Results of a statistical treatment of the first level's lexic of the distributive dictionary /DD/ of Russian and a comparison of the data of Russian and French DD]. Vestnik Xar'kovskogo universiteta 1979/185. Inostrannyj jazyk, vyp. 12, 7-11.
31. GEENS, D.: On measurement of lexical differences by means

- of frequency. Altmann, G. (Ed.): Glottometrika 1, Bochum: Brockmeyer, 1978, 46-72.
32. KAASIK, U., TULDAVA, J., ÄÄREMAA, K.: Eesti keele sõnavormide pöördagedussõnastik [Reverse frequency dictionary of word forms in Estonian]. Soontak, Ja. (Ed.): Eesti keele sõnavarastatistika eriküsimusi [Special Problems of Lexicostatistics of Estonian]. Tartu: State University of Tartu, 1980, 5-153.
33. KIJAK, T.R., KOSTAŠ, Ju.P., SKOROXOD'KO, E.F.: Ocenka stepeni motivirovannosti nemeckoj i russkoj terminologičeskoj i obščepotrebitel'noj leksiki [Estimation of the degree of motivation of Russian and German terminological and general lexicon]. Strukturnaja i matematičeskaja lingvistika 8, 1980, 37-44.
34. NESTEROV, M.N.: Častotnye slovari russkogo jazyka [Frequency dictionaries of the Russian language]. Russkaja reč' 1979/3, 108-112.
35. SOONTAK, Ja. (Ed.): Eesti keele sõnavarastatistika eriküsimusi [Special Problems of Lexicostatistics of Estonian]. Tartu: State University of Tartu, 1980, 168 pp.
36. VANNIKOV, Ju.V.: Edinaja spravočnaja lingvo-statističeskaja sistema kak baza učebnoj leksikografii: Stat'ja vtoraja [A homogeneous linguo-statistical reference system as a basis of school-lexicography: Second article]. Pervodnaja i učebnaja leksikografija. Moskva 1979, 124-148.
37. VESLER, I.Ju.: Algoritm sžatija russkix slov, učityvajuščij dlinu slova [A word-length-sensitive algorithm for the compression of Russian words]. Strukturnaja i matematičeskaja lingvistika 8, 1980, 14-22.

TEXT ANALYSIS

38. ALTMANN, G.: Zur Anwendung der Quotiente in der Textanalyse. Altmann, G. (Ed.): Glottometrika 1, Bochum: Brockmeyer,

- 1978, 91-106.
39. BRAINERD, B.: Pronouns and genre in Shakespeare's drama. Computers and the Humanities 13, 1979, 3-16.
 40. Častotnyj slovar' romana L.N. Tolstogo "Vojna i mir" [A Frequency Dictionary of L.N. Tolstoy's Novel "War and Peace"]. Tula: Gosudarstvennyj pedagogičeskij institut, 1978, 378 pp.
 41. HARMATIUK, S.J.: A statistical approach to some aspects of style in six old English poems: a computer assisted study. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer 1980, 89-107.
 42. KLEINLOGEL, A.: Fundamentals of a formal theory of manuscript classification and genealogy. Colloques Internationaux du C.N.R.S. No. 579. La pratique des ordinateurs dans la critique des textes 1980, 193-205.
 43. KUDRJAVCEVA, T.A., VANNIKOV, Ju.V., ABDALJAN, I.P.: Universal'nyj karakter formuly čitabel'nosti Ju.A. Tuldavy i vozmožnost' ee izpol'zovanija v kačestve étalona [The universal character of Y.A. Tuldava's readability formula and the possibility of its use as a standard]. Lingvistica 11, 1979, 50-60.
 44. LEBEDEV, B.M.: Lingvostatističeskoe modelirovanie nemeckogo naučno-texničeskogo teksta: Pod-jazyk televizionnoj tekhniki [The Linguo-statistical Modelling of a German Scientific-technical Text: The Sublanguage of Television Engineering]. Aftoreferat kandidatskoj dissertacii. Leningrad 1979, 19 pp.
 45. McKINNON, A.: Most frequent words and the clustering of Kierkegaard's works. Style 12/1978, 241-257.
 46. McKINNON, A.: Aberrant frequency words: their identification and uses. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer 1980, 108-124.
 47. MOSKOWICH, W., CAPLAN, R.: Distributive-statistical text analysis: A new tool for semantic and stylistic research.

- Altmann, G. (Ed.): Glottometrika 1, Bochum: Brockmeyer, 1978, 107-153.
47. NADAREJŠVILI, I.S.: Statističeskij analiz morfoložičeskoj struktury teksta (na materiale prozy K. Gamsuxurdija) [Statistical analysis of the morphological structure of texts (on material of K. Gamsuxurdija's prose)]. Lingvistica 11, 1979, 80-87.
 48. NOWAKOWSKA, M.: Some formal aspects of dynamics of dialogues. Altmann, G. (Ed.): Glottometrika 1, Bochum: Brockmeyer, 1978, 73-90.
 49. PIEPER, U.: Über die Aussagekraft statistischer Methoden für die linguistische Stilanalyse. Ars Linguistica 5, Tübingen: Gunter Narr, 1979, 157 pp.
 50. PRIKOD'KO, S.M., SLAST'EN, T.P.: K voprosu o raspredelenii mežfrazovyx svjazej v tekste [On the question of the distribution of interphrasal connections in texts]. Strukturnaja i matematičeskaja lingvistika 8, 1980, 93-98.
 51. RATKOWSKY, D.A., HALSTEAD, M.H., HANTRAIS, L.: Measuring vocabulary richness in literary works: A new proposal and a re-assessment of some earlier measures. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer 1980, 125-147.
 52. SADYKOV, T.: K issledovaniju kvantitativno-lingvističeskoj struktury kirgizskogo teksta na ÉVM: O častotnom slovar'e eposa "Manas" [Towards an investigation of the quantitative-linguistic structure of a Kirghiz text by means of a computer: On the frequency dictionary of the epos "Manas"]. Izvestija Akademii nauk Kirgizskoj SSR, Frunze 1979/3, 93-100.
 53. ŠAJKEVIČ, A.Ja.: Differenciacija statističeskix klassifikacij tekstov [Differentiation of statistical classifications of texts]. Lingvistica 11, 1979, 100-106.
 54. ŠULJAK, N.T., MAKARENKO, V.G.: Opredelenie ob-ema vyborki dlja statističeskogo analiza sintaksisa teksta (na mate-

riale ukraïnskogo jazyka) [The determination of sample size for the statistical analysis of text syntax (based on Ukrainian)]. Strukturnaja i matematičeskaja lingvistika 8, 1980, 111-115.

55. STERNER, H.: Wahrscheinlichkeitstheoretisches Modell für das Zusammentreffen der Personen im Drama. Grotjahn, R. (Ed.): Glottometrika 2, Bochum: Brockmeyer 1980, 148-157.
56. TULDAVA, Ju.: O nekotoryx kvantitativno-sistemnyx xarakteristikax polisemii [On some quantitative-systemic features of polysemy]. Linguistica 11, 1979, 107-139.
57. TUTTLE, H.: Clues to vocabularic structure: Comparative suffixal productivity in the five books of Rabelais. Lingua e Stile 14, 1979, 151-163.

PSYCHOLINGUISTICS

58. ADAMS, M.J.: Models of word recognition. Cognitive Psychology 11, 1979, 133-176.
59. AKIYAMA, M.M., BREWER, W.F., SHOBEN, E.J.: The yes-no question answering system and statement verification. Journal of Verbal Learning and Verbal Behavior 18, 1979, 365-380.
60. BERRIAN, R.W., METZLER, D.P., KROLL, N.E.A., CLARK-MEYERS, G.M.: Estimates of imagery, ease of definition and animateness for 328 adjectives. Journal of Experimental Psychology, Human Learning and Memory 5, 1979, 435-447.
61. BIEDERMAN, J., YAO-CHUNG TSAO: On processing Chinese ideographs and English words: Some implications from Stroop test results. Cognitive Psychology 11, 1979, 125-132.
62. BYRNE, B.: Rules of prenominal adjective order and the interpretation of "incompatible" adjective pairs. Journal of Verbal Learning and Verbal Behavior 18, 1979, 73-78.
63. CARON, J.: La compréhension d'un connecteur polysémique:

La conjonction "si". Bulletin de Psychologie 32, 1978/1979, 791-801.

64. CHAMBERS, S.M.: Letter and order information in lexical access. Journal of Verbal Learning and Verbal Behavior 18, 1979, 225-241.
65. CLARK, H.H.: Responding to indirect speech acts. Cognitive Psychology 11, 1979, 430-477.
66. CLARK, H.H., SENGUL, C.J.: In search of referents for nouns and pronouns. Memory and Cognition 7, 1979, 35-41.
67. COLEMAN, E.B.: Generalization effects vs. random effects: Is σ^2_{TL} a source of type 1 or type 2 error? Journal of Verbal Learning and Verbal Behavior 18, 1979, 243-256.
68. CUTLER, A., FODOR, J.A.: Semantic focus and sentence comprehension. Cognition 7, 1979, 49-59.
69. FRAUENFELDER, U., DOMMERGUES, J.Y., MEHLER, J., SEGUI, J.: L'intégration perceptive des phrases. Bulletin de Psychologie 32, 1978/1979, 893-902.
70. GLANZER, M., EHRENREICH, S.L.: Structure and search of the internal lexicon. Journal of Verbal Learning and Verbal Behavior 18, 1979, 381-398.
71. GREEN, T.R.G.: The necessity of syntax markers: Two experiments with artificial languages. Journal of Verbal Learning and Verbal Behavior 18, 1979, 481-496.
72. GROSJEAN, F., GROSJEAN, L., LANE, H.: The patterns of silence: Performance structures in sentence production. Cognitive Psychology 11, 1979, 58-81.
73. HAMPTON, J.A.: Polymorphous concepts in semantic memory. Journal of Verbal Learning and Verbal Behavior 18, 1979, 441-461.
74. HAYES-ROTH, B., THORNDYKE, P.W.: Integration of knowledge from text. Journal of Verbal Learning and Verbal Behavior 18, 1979, 91-108.
75. HUTTENLOCHER, J., LUI, F.: The semantic organization of so-

- me simple nouns and verbs. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 141-162.
76. KASTNER, M., SCHMACHTENBERG, C.: Eine experimentelle Untersuchung zur semantischen Integriertheit von Texten. *Psychologische Beiträge* 21, 1979, 78-97.
 77. KING, M.C., BEVAN, W.: Pictures, words, and the structure of the trace in immediate recall. *Bulletin of the Psychonomic Society* 14, 1979, 155-157.
 78. LESGOLD, A.M., ROTH, S.F., CURTIS, M.E.: Foregrounding effects in discourse comprehension. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 291-308.
 79. MÄGISTE, E.: The competing language systems of the multilingual: A developmental study of decoding and encoding processes. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 79-89.
 80. MARX, W.: Die Dominanz des Substantivs als Träger der assoziativen Bedeutung. *Zeitschrift für experimentelle und angewandte Psychologie* 26, 1979, 596-602.
 81. MARX, W., SCHÜRER-NECKER, E.: Überlegungen zur Interpretation des Zipfschen Gesetzes am Beispiel der frühen Kindersprache. Altmann, G. (Ed.): *Glottometrika* 1, Bochum: Brockmeyer, 1978, 154-167.
 82. MOORE, T.E., BIEDERMAN, J.: Speeded recognition of ungrammaticality: Double violations. *Cognition* 7, 1979, 285-299.
 83. PINKER, S., BIRDSOING, D.: Speaker's sensitivity to rules of frozen word order. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 497-508.
 84. POTTER, M.C., FAULCONER, B.A.: Understanding noun phrases. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 509-521.
 85. RAEURN, V.P.: The role of the verb in sentence memory. *Memory and Cognition* 7, 1979, 133-140.

86. REMEZ, R.E.: Adaptation of the category boundary between speech and nonspeech: A case against feature detectors. *Cognitive Psychology* 11, 1979, 38-57.
87. SCHULTZ, E.E.jr., KAMIL, M.L.: The role of some definite references in linguistic integration. *Memory and Cognition* 7, 1979, 42-49.
88. SHATTUCK-HUFNAGEL, S., KLATT, D.H.: The limited use of distinctive features and markedness in speech production: Evidence from speech error data. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 41-55.
89. SOLMECKE, G., BOOSCH, A.: Entwicklung eines Eindrucksdifferentials zur Erfassung von Einstellungen gegenüber Sprachen. *Linguistische Berichte* 60, 1979, 46-64.
90. STANNERS, R.F., NEISER, J.J., HERNON, W.P., HALL, R.: Memory representation for morphologically related words. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 399-412.
91. STRUBE, G.: Associative meaning and its analysis. Grotjahn, R. (ed.): *Glottometrika* 2, Bochum: Brockmeyer 1980, 158-197.
92. SWINNEY, D.A., CUTLER, A.: The access and processing of idiomatic expressions. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 523-534.
93. SWINNEY, D.A., ONIFER, W., PRATHER, P., HIRSHKOWITZ, M.: Semantic facilitation across sensory modalities in the processing of individual words and sentences. *Memory and Cognition* 7, 1979, 159-165.
94. TAFT, M.: Recognition of affixed words and the word frequency effect. *Memory and Cognition* 7, 1979, 263-272.
95. TAFT, M.: Lexical access via an orthographic code: The basic orthographic syllabic structure (BOSS). *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 21-39.
96. TANNENHAUS, M.K., LEIMAN, J.M., SEIDENBERG, M.S.: Evidence

for multiple stages in the processing of ambiguous words in syntactic contexts. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 427-440.

97. Van der MOLEN, H., MORTON, J.: Remembering plurals: Unit of coding and form of coding during serial recall. *Cognition* 7, 1979, 35-47.

98. WENDER, K.F., GLOWALLA, U.: Models for within-proposition representation tested by cued recall. *Memory and Cognition* 7, 1979, 401-409.

99. WHITTEN, W.B. II., NEWTON SUTER, W., FRANK, M.L.: Bidirectional synonym ratings of 464 noun pairs. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 109-127.

100. YEKOVICH, F.R., WALKER, C.H., BLACKMAN, H.S.: The role of presupposed and focal information in integrating sentences. *Journal of Verbal Learning and Verbal Behavior* 18, 1979, 535-548.

S O C I O L I N G U I S T I C S

101. RODRIGUE-SCHWARZWALD, O.: The influence of demographic variables on linguistic performance. Altmann, G. (Ed.): *Glottometrika* 1, Bochum: Brockmeyer, 1978, 173-197.

H I S T O R I C A L L I N G U I S T I C S

102. GUITER, H.: A propos de glottochronologie. *Cahiers de l'Institut Linguistique de Louvain* 4, 1977/3-4, 3-19.

103. GUITER, H.: Encore la glottochronologie. *Cahiers de l'Institut Linguistique de Louvain* 4, 1977/3-4, 3-28.

104. GUITER, H.: Glottochronologie et langues occidentales. *Cahiers de l'Institut Linguistique de Louvain* 4, 1977/2, 1-32.

105. KUPČINSKIJ, O.A.: Statistika i stratigrafija vostočno-slavjanskix toponimov na -IČI i nekotorye voprosy isto-ričeskoj geografii zaselenija: na materiale USSR [Statistics and stratigraphy of East Slavonian toponyms ending on -iči and some questions of historical settlement-geography: on the material of the Ukrainian Soviet Socialistic Republic]. *Voprosy geografii* 110, 1979, 89-103.

106. MAŃCZAK, W.: Les lois du développement analogique. *Linguistics* 205, 1978, 53-60.

107. MAŃCZAK, W.: Frequenz und Sprachwandel. Lüdthke, H. (Ed.): *Kommunikationstheoretische Grundlagen des Sprachwandels*. Berlin: de Gruyter, 1980, 37-79.

108. PALAGINA, V.V.: Častotnyj slovar' tamožennyx knig Tomska 1624-1627 gg. [The frequency dictionary of the customs-books of Tomsk 1624-1627]. *Russkie govory v Sibiri*. Tomsk 1979, 78-83.

B I B L I O G R A P H Y

109. Current Bibliography. Grotjahn, R. (Ed.): *Glottometrika* 2, Bochum: Brockmeyer, 1980, 198-218.

110. Current Bibliography. Altmann, G. (Ed.): *Glottometrika* 1, Bochum: Brockmeyer, 1978, 220-231.

111. KAMINSKA, I.: Bibliography of quantitative linguistics in Poland. Altmann, G. (Ed.): *Glottometrika* 1, Bochum: Brockmeyer, 1978, 198-219.

D O C U M E N T A T I O N

112. LANGENDÖRFER, H., HOFFMANN, F.A.: Entwurf und Implementierung eines Literatur-Dokumentationssystems mit direktem Zugriff. *Sprache und Datenverarbeitung* 2, 1978, 153-158.

DI A L E C T O L O G Y

113. GUITER, H.: Français Central et dialects du Nord-Ouest selon l'Atlas linguistique de la France. Bulletin Philologique et Historique 1976, 55-66.
114. PŠENIČNOVA, N.N.: Mera specifičnosti i nekotorye voprosy klassifikacii častnyx dialektnyx sistem [A measure of specificity and some questions of the classification of single dialectal systems]. Obščeslavjanskij lingvističeskij atlas. Materialy i issledovanija. 1977. Moskva: Nauka, 1979, 236-249.
115. STELLMACHER, D.: Studien zur gesprochenen Sprache in Niedersachsen. Eine soziolinguistische Untersuchung. Marburg: Elwert, 1977, 224 pp.

L A N G U A G E T E A C H I N G

116. Škola-seminar po optimizaciji prepodovanija inostrannyx jazykov s pomošč'ju tehničeskix sredstv (tezisy dokladov i soobščenijs) [A School-Seminar About the Optimization of Foreign-Language Tuition by Means of Technical Media (Thesises of Lectures and Reports)]. Kišinev 1979, 129 pp.

All contributions to the CURRENT BIBLIOGRAPHY should be sent to Prof.Dr. Werner Lehfeldt, Universität Konstanz, Philosophische Fakultät, Fachgruppe Sprachwissenschaft Slavistik, P.O. Box 5560, D 7750 Konstanz.

List of contributors

- Altmann, Gabriel, Ruhr-Universität Bochum, Sprachwissenschaftliches Institut, Postfach 102148, 4630 Bochum 1
- Boroda, Moisei G., Konservatorija, ul.Griboedova, 380004 Tbilisi 4, USSR
- Fischer, Erik, Ruhr-Universität Bochum, Germanistisches Institut, Postfach 102148, 4630 Bochum 1, West Germany
- Grotjahn, Rüdiger, Ruhr-Universität Bochum, Seminar für Sprachlehrforschung, Postfach 102148, 4630 Bochum 1, West Germany
- Kolman, Ludek, Institut rizeni, Jungmannova 29, 11549 Praha 1, Czechoslovakia
- Marx, Wolfgang, Universität München, Psychologisches Institut, Geschwister-Scholl-Platz 1, 8000 München, West Germany
- Matthäus, Wolfhart, Ruhr-Universität Bochum, Psychologisches Institut, Postfach 102148, 4630 Bochum 1, West Germany
- Rehak, Jan, Vinohradská 67, 12000 Praha 2, Czechoslovakia
- Rehakova, Blanka, Vinohradska 67, 12000 Praha 2, Czechoslovakia
- Schäfer, W.-D., Dorstener Str.90, 4630 Bochum, West Germany
- Stoffer, Thomas, Ruhr-Universität Bochum, Psychologisches Institut, Postfach 102148, Bochum 1, West Germany
- Strube, G., Universität München, Psychologisches Institut, Geschwister-Scholl-Platz 1, 8000 München, West Germany
- Wildgruber J., Dorstener Str.90, 4630 Bochum, West Germany

INSTRUCTIONS TO AUTHORS

1. Contributions should be in English, German or French, but in English if at all possible.
2. MSS should be sent in duplicate to G.Altmann, Sprachwissenschaftliches Institut der Ruhr-Universität Bochum, Postfach 102148, D-4600 Bochum 1, West Germany, or to a member of the editorial board.
3. An abstract in English should be included at the beginning of the contribution.
4. Brief information about the author(s) (e.g. institution, position, main research interests) should be included on a separate page.
5. Authors are responsible for obtaining permission to publish material (figures, tables, etc.) for which they do not hold the copyright.
6. Each contributor to Glottometrika will receive one copy of the volume free of charge. Offprints may be purchased from the publisher if ordered at the proof stage.
7. Preparation of the manuscript. In order to facilitate the editing of the series the MSS should be submitted in final form, ready for photo-offset, observing the conventions delineated below. If this is not possible the MSS will be retyped at the expense of the editors.

Conventions:

Title: IBM Selectric typeface Orator if possible; 4 cm from top, centered.

Name and institutional affiliation (+ place if necessary): IBM Courier Italic (if possible), 2 cm under title.

Abstract: single-spaced, maximum of 20 lines, IBM Courier 10, 2 cm below name and institution.

Text: Courier 10, 1 1/2-spaced, 2 cm below the abstract should come the first line.

Paragraphs: indent 3 spaces.

Margins: 2 cm left as well as right, 3 cm above and below.

Chapter or paragraph titles, subtitles: any contrasting typeface.

Emphasized words within text as well as italics (i.e. for foreign words): underlined or in italics.

Notes and Bibliography should be listed separately at the end of the article.

The Bibliography should list the literature used in alphabetical order by author's last name, in the following form:

Brainerd, B., Graphs, topology and text. Poetics, 1977, 6, 1-14
Charpentier, C.M., La statistique linguistique pratiquée dans les pays anglo-saxons. Éléments de bibliographie (1970-1976). In: David, J. & Martin, R. (Eds.), Études de statistique linguistique. Paris: Klincksieck, 1977, 89-119

Titles of journals should not be abbreviated.

References both in the text and in the notes should take the form of the following examples: (cf. Brainerd, 1977:10); (David & Martin, 1977:12).

Figures and graphs should be ready for offset printing and be large enough to remain legible after a reduction in size of 50%. Furthermore they should be numbered consecutively, identified at the bottom with legends and, if possible, placed at an appropriate spot in the running text.

Tables should be numbered consecutively, provided with titles at the top and, if necessary, with legends (also at the top) and should be typed single-spaced. If at all possible, tables should be integrated into the running text, just as with figures and graphs.