

QUANTITATIVE LINGUISTICS

Vol. 14

GLOTTOMETRIKA 4

edited

by

W. Lehfeldt

U. Strauss



Studienverlag Dr. N. Brockmeyer
Bochum 1982

QUANTITATIVE LINGUISTICS

Editors

G. Altmann, Bochum
R. Grotjahn, Bochum

Editorial Board

N. D. Andreev, Leningrad
M. V. Arapov, Moscow
B. Brainerd, Toronto
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
4. Ed. by W. Lehfeldt; U. Strauss. – 1982.
(Quantitative linguistics; Vol. 14)
ISBN 3-88339-250-2

NE: Lehfeldt, Werner [Hrsg.]; GT

ISBN 3-88339-250-2

Alle Rechte vorbehalten

© 1982 by Studienverlag Dr. N. Brockmeyer

Querenburger Höhe 281, 4630 Bochum 1

Druck Thiebes GmbH & Co Kommanditgesellschaft Hagen

We would like to express our gratitude
to the

STIFTUNG VOLKSWAGENWERK

a generous grant from which made possible the
publication of the German edition of this book.

INHALT

WINKLER, P. Quantitative Analyse phonetischer Transkripte (Umgangssprache)	1
DIETZE, J. Grapheme und Graphemkombinatorik der russischen Fachsprache. Eine phonostatistische Untersuchung	80
GERLACH, R. Zur Überprüfung des Menzerath'schen Gesetzes im Bereich der Morphologie	95
KÖHLER, R. Das Menzerathsche Gesetz auf Satzebene	103
CHISHOLM, David H. Phonological Patterning in English and German Verse: A Computerassisted Approach	114
ALTMANN, G., BURDINSKI, V. Towards a Law of Word Repetitions in Text-Blocks	147
LAUBER, J., LINHART, J. Kanonische Analyse von Kontingenztafeln	168
CURRENT BIBLIOGRAPHY	183

QUANTITATIVE ANALYSE PHONETISCHER TRANSKRIPTE (UMGANGSSPRACHE)

Peter Winkler
Universität Konstanz

1.1 EINLEITUNG

Statistische Analysen von Lauthäufigkeiten und Lautkombinationen sind in der Linguistik und Phonetik schon verschiedentlich durchgeführt worden. SEGAL (1972, 22ff.) zitiert eine Reihe von Autoren, die bereits vor der Jahrhundertwende statistische Auswertungen von Sprachmaterial publiziert haben, das entweder schrifttextlich oder in phonetischen Transkriptionen vorlag. Später hatte die Phonologie die Anwendung quantitativer Methoden in der Linguistik besonders gefördert und angeregt; die Mehrzahl statistischer Sprachlautanalysen sind phonematisch orientiert. Mit phonetischer Zielsetzung sind noch weitere Applikationen quantitativ-statistischer Methoden möglich, von denen einige im Überblick von KARLGREN (1974³, 129ff.) aufgeführt sind. Auch hier ließe sich eine Tradition bis vor 1900 zurückverfolgen, in der in jüngerer Zeit die Phonometrie von ZWIRNER das prominenteste Beispiel ist.

Quantitative Auswertungen phonetischer Transkripte sind im Wesentlichen für drei verschiedene Bereiche eingesetzt worden:

(1) Methodologie: Untersuchung der phänomenalen und perzeptiven Leistung von Transkribenten (z.B. OLLER und EILERS 1975 über phonetische Erwartbarkeit); Leistungsspektren von Transkriptionssystemen, Validitäts- und Reliabilitätsstudien (vgl. RICHTER und RICHTER 1981).

(2) Ausspracheuntersuchungen: Standard-, Umgangs- und Dialektaussprache, alltägliche, künstlerische, rhetorische Artikulationsstile; Basisuntersuchungen für die Kodifizierung von Lautungsformstufen oder für die Herstellung von Übungsmaterial (z.B. ORTMANN 1975), Sprechbewegungsabläufe (LINDNER 1975), soziolinguistische Erhebungen zur Aussprachevariation u.a.

(3) Paralinguistik: Beschreibung sprechereigener Artikulationsbesonderheiten, situations- und kontextgebundener Unterschiede, der Anpassung verschiedener Sprecher untereinander im Artikulationsniveau usw.

In der nachfolgenden quantitativen Analyse wird je ein Einzelproblem aus den o.g. Bereichen bearbeitet:

(1) Die auditive Festlegung von Lautlängen (Segmentdauer) bei ohrenphonetischen Segmentationen.

(2) Lauthäufigkeit und -kombinatorik umgangssprachlich beeinflusster Aussprache.

(3) Individuelle Artikulationseigenheiten zweier Sprecher während eines dyadischen Gesprächs im colloquial style.

1.2 MATERIAL UND METHODEN

Es werden die Transkripte der Tonaufnahmen zu TAKE D des Konstanzer Forschungsprojektes "Analyse unmittelbarer Kommunikation und Interaktion als Zugang zum Problem der Entstehung sozialwissenschaftlicher Daten" (finanziert von der Fritz-Thyssen-Stiftung; Leitung: Prof. Th. Luckmann/ Prof. Dr. P. Gross) zugrundegelegt.

(1) Sprachmaterial: Die Aufnahmen entstammen einer realen Gesprächssituation, in der untrainierte Sprecher (ein weiblicher und ein männlicher Gesprächsteilnehmer) frei ein zehnminütiges Gespräch gestalteten. Das Rahmenthema dieser Dyade wurde nur global vorgegeben. Die Aufnahmen enthalten "small-Talk"-Sequenzen, die in umgangssprachlich beeinflussten Sprechstil ausgeführt wurden. Die Sprecherin läßt in der Artikulation, teilweise auch in der Wortwahl schwäbisch-alemannische Dialektanklänge erkennen; der Sprecher zeigt Einflüsse der berliner Aussprache. Typisch für die Umgangsaussprache ist auch in dieser Aufnahme, daß die Sprecher frei zwischen hochgelauteten und dialektal eingefärbten Realisierungen oszillieren. Der schwäbisch-alemannische Einfluß zeigte sich u.a. im Ersatz von $[-st]$ durch $[-t]$, in der Reduktion von $[a]$ im Wort "das" zu $[ə]$ und in Vokalrundungen bzw. im Ersatz von $[a]$ durch $[ɔ]$. Als Merkmale der berliner Aussprache traten z.B. übermäßige Öffnungen der vokalisch aufgelösten R-Laute der Kombination $-er$ und Ersatzrealisationen von $[ts]$ durch $[s]$ auf. Bei beiden Sprechern waren die üblichen Verschleifungen, Elisionen und Neubildungen typisch für den colloquial style (z.B. wurde der Satzanfang "Na, weil ich es..." als $[navəks]$ ausgesprochen).

(2) Aufnahme: Normalhalliger Raum (Filmstudio); Tonaufnahme mit für jeden Sprecher separaten Mikrofonen und Aufnahmegeräten (1 Pärchen Beyer Dynamic Mikrofone mit Richtwirkung; 2 NAGRA-Tonbandmaschinen mit 19 cm/sec; Abstand der Mikrofone vom Sprecher ca. 1m).

(3) Segmentation und Transkription: Die Tonaufnahmen wurden einzeln digitalisiert und jeweils gesondert computerunterstützt segmentiert sowie in broad transcription nach dem IPA-System transkribiert (unter Benutzung der PDP-11/50 Rechenanlage und der Software des Instituts für Phonetik und sprachliche Kommunikation der Universität München; Leitung: Prof. Dr. H.G. Tillmann; Beschreibung des Segmentations- und Transkriptionssystems vgl. WINKLER 1981, 28 ff.).

2.1 AUDITIVE FESTLEGUNG VON LAUTEXTENSIONEN (LÄNGE VON SEGMENTEN)

TIFFANY und CARRELL (1977², 45 ff.) bezeichnen die Transkriptionsmethode als eine perzeptuelle Analyse der "sound structure of language and articulatory possibilities", wobei sie vorsichtig einschränken, daß "...phonetic transcription is a record of the listener's impression of what he or she hears." (60). Das entspricht einer weitverbreiteten methodologischen Auffassung über die Transkriptionsmethode, in deren Rahmen bestimmte "Mängel" der Methode im Wesentlichen auf subjektive und individuelle Perzeptionen, Abweichungen von der interindividuellen Norm usw. zurückgeführt werden. Demgegenüber sieht die Phänomenologie in den immer wieder auftretenden Diskrepanzen zwischen den Notaten verschiedener Transkribenten oder zwischen Schallstruktur und Transkriptionen keine subjektiven Einzelfälle, sondern generelle Strukturen der menschlichen Perzeptions- und Apperzeptionsvorgänge. Verkürzt ausgedrückt beschreibt ein Transkript nicht die Schallstruktur von Lauten oder die Artikulationsmöglichkeiten der Sprechwerkzeuge, sondern die apperzeptiven Möglichkeiten zum Erkennen von Artikulationsbewegungen oder Sprachzeichen bzw. deren Varianten. In der Konsequenz der phänomenologischen Position liegt der Verzicht auf den Anspruch einer sogenannten "Abbildungstreue" von Notaten und weiterhin auch auf die Möglichkeit, mit ohrenphonetischen Methoden den Sprachschall oder die Sprechphysiologie so untersuchen zu können, wie es die signal-

und artikulationsphonetischen Methoden tun. Die phonetische Transkription kann kein Ersatz für Schall- und Lautanalysen sein, sondern sie dokumentiert den phänomenal-apperzeptiven Bereich gesprochener Sprache.

Analysen phonetischer Transkripte sind, immer der phänomenologischen Position zufolge, Untersuchungen zur Struktur der intersubjektiven Konstitution von Sprachzeichen aus Prototypen von Schallkonfigurationen (Phonemsystem), zur Technik des analytischen und funktionellen Hörens, zum apperzeptiven Aufbau von phonetischen Idealtypisierungen (Lauten), zur natürlichen Fähigkeit des Menschen, aus Schallprodukten auf die sprechphysiologischen Bewegungen schließen zu können, und zur fachspezifischen Wissensstruktur, den spezifischen Wissensvorräten und Elementen des Vorwissens. Die Phonologie kann zurecht Transkriptionen für die Untersuchung des Phonemsystems heranziehen, ebenso wie die Phonetik Notationen zur Beschreibung der diskriminierbaren Laute, Lautvarianten, Suprasegmentale und der erschließbaren Organbewegungen verwendet, nicht aber zur Analyse von Vorgängen, die während der normalen Kommunikation transphänomenal bleiben (z.B. die akustische Schallgestalt oder die tatsächlichen Sprechbewegungen).

Schon die Aufgliederung des Schallkontinuums in Laute (Segmentierung) ist eine rein phänomenal begründete Prozedur, bei der diskrete Muster von Lauttypen und Untertypen auf den Schallstrom "aufgeprägt" werden. Die akustische Struktur des Schalls ist als ein "Signal" aufzufassen, das durch weitreichende Apperzeptionsprozesse zur "Nachricht" wird (so beschreibt es das Signal-Nachricht-Modell der modernen Phonetik). Bei der computerunterstützten Segmentation, bei der der Schall gleichzeitig gesehen und gehört werden kann, werden die Unterschiede zwischen Signalstruktur und auditivem Korrelat unmittelbar evident.

Zum Beispiel kann sich während der Segmentationsarbeit der Transkribent entweder mehr nach dem visualisierten Schall (Oszillogramm) richten oder mehr nach dem auditiven Eindruck. Die Grenzen der Segmente würden jeweils anders festgelegt werden.

Als Beispiel sind in Abb. 1 das segmentierte Oszillogramm und fünf weitere Parameter des Anfangs von TAKE D, Sprecherin, angeführt. Die Diskrepanz von Schallverlauf und Segmenteinteil-

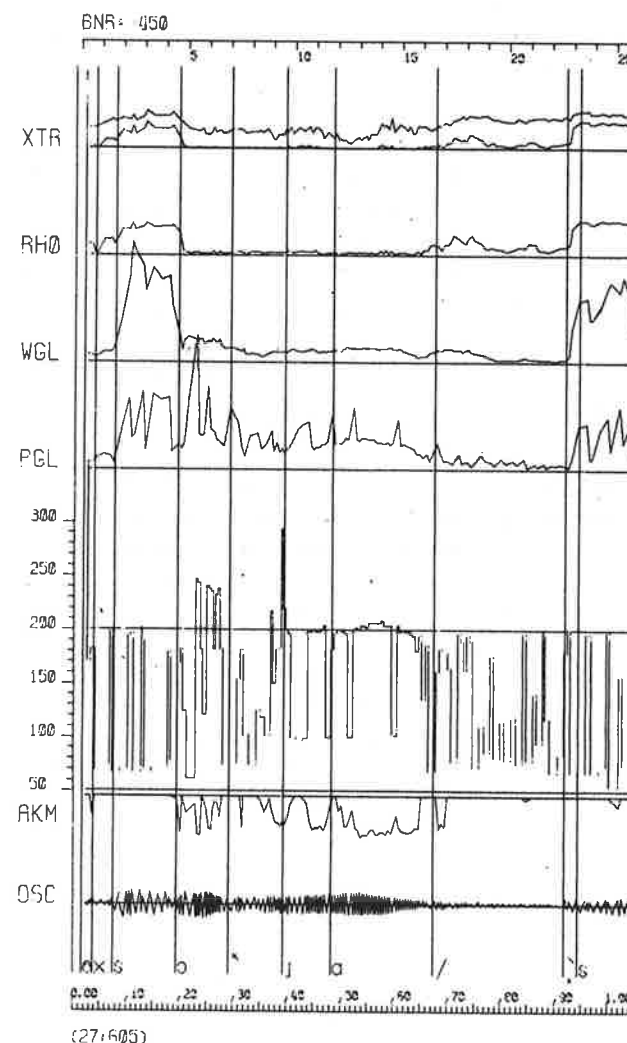


Abb.1: Segmentiertes Oszillogramm und Parameter-Darstellung des Anfangs von TAKE D, Sprecherin

lung aufgrund auditiver Kriterien ist deutlich bei der Kombination von [s] und [o] zu sehen. Die akustische Einteilung würde den Laut [o] später anfangen lassen; etwa mit dem Einsetzen des Stimmtones (Kurve über "AKM" in treppenförmiger Abstufung bei ca. 240 Hz) und dem gleichzeitigen Pegelanstieg (Kurve PGL).

Die auditive Bestimmung von Lautextensionen ist zwar nicht völlig unabhängig vom Schallverlauf, sie folgt aber zugleich einer eigenen, phänomenalen Regelhaftigkeit. Sie hat so gut wie keinen Bezug zur Umschrift der phonematischen Länge oder Kürze eines Lautes: Als "lang" transkribierte Vokale können so kurz sein wie "kurze" Laute und umgekehrt. (In der Phonetik wird diese Erscheinung damit erklärt, daß die Vokalqualität in der deutschen Sprache die -quantität dominiert: Offene Vokale werden zugleich "kurz" gehört, geschlossene "lang".) Die Transkription der (phonematischen) Lautdauer ist tatsächlich keine perzeptuelle Analyse der Lautextension, sondern die Kennzeichnung von bedeutungsunterscheidender Dauer eines Sprachzeichens. Die Konstitution von bedeutungsdifferenzierenden Elementen ist nur mittelbar von der akustischen Substanz abhängig.

Dazu zwei Beispiele aus dem untersuchten Material: Die Zeitstrecke zwischen Segmentanfang und -ende (Extension des Lautes) wurde bei einer Samplingfrequenz von 20KHz mit der Genauigkeit von 1/20 000 sec. digital bestimmt und in sec. umgerechnet. Die auditiv festgesetzten Zeitspannen für alle Realisationen ein- und desselben Phonems wurden maschinell sortiert und die statistischen Parameter $\bar{x}$, s und s^2 berechnet.

Tabelle 1: Auditive Längen der Vokale [ø], [œ], [o] und [ɔ]
TAKE D, Sprecherin, ca. 10 Min. Gespräch

	[ø]	[œ]	[o]	[ɔ]
sec	.22645	.2304	.2177	.2325
	.2038	.2304	.20845	.2291
	.16875	.1408	.20845	.2176
	.1612	.1408	.20665	.2105
	.1048	.12295	.20505	.2038
	.1036	.12125	.2044	.192
	.1002	.1024	.1792	.1916
	.0512	.09575	.1792	.18835
		.0768	.1696	.1786
		.0767	.1664	.16875
		.07615	...	...
		.07175	...	...
		.064	.0384	.0384
		.0595	.0384	.03655
		.05785	.0377	.0362
		.0512	.0336	.0356
		.0434	.03335	.03385
		.0403	.026	.0303
			.0064	.0256
			.0064	.0256
			.0064	.0256
			.0064	.0256
			.0064	.0256
n	8	19	86	88
$\bar{x}$	1.1199	1.872	8.1288	8.3971
s^2	.13999	.09891	.09452	.09542
s	.00355	.00306	.00272	.00277
s	.05959	.05360	.05216	.05266

Die Absolutzeitwerte sind bei [o] niedriger als bei [ɔ] (Extremzeitwerte), wobei im mittleren Bereich gleiche Zeitwerte auftreten (z.B. .0384 sec.). Beim geschlossenen und offenen OE-Laut sind die Extremwerte ähnlich; es treten wiederum identische Längen auf (.0512 sec.).

Die theoretisch zu erwartenden Minimallängen - das Minimum, das beim Segmentieren nicht unterschritten werden kann, weil der Laut dann in eine andere Lautklasse eingeordnet wird - kann annähernd durch lineare Regression ($x = 0$) geschätzt werden.

	Minimum	empirisch	r (Anpassungsgüte)
[ø]	.033	.0512	.97
[æ]	.015	.0403	.87
[o]	.0003	.0064	.98
[ɔ]	.0213	.0256	.95

Die durchschnittliche Länge ist bei [ɔ] etwas größer als bei [o]; [ø] ist im Durchschnitt länger als [æ]. Zum Vergleich der Standardabweichungen wird zunächst der Variabilitätskoeffizient berechnet: $V = 42.57[\phi], = 54.19[\alpha], = 55.28[o], = 55.19[\sigma]$. Mit Ausnahme von [ø] liegen die Variationsbreiten in ähnlichen Größenordnungen.

Die Beispiele zeigten, in welchen Zeitdimensionen Segmente festgelegt werden. In der Tabelle 2 sind die Maximal- und Minimaextensionen aller Segmente des Materials aufgelistet. Die zwischen den Extremwerten liegenden Quantelungen lassen sich theoretisch berechnen: Die auditiven Zeitabschnitte steigen vom kürzesten bis längsten Segment eines Lauttyps logarithmisch (ln) oder wie 3. Wurzel aus einem Polynom dritten Grades an. Beim Segmentieren wird demnach ein "längerer" Laut nicht doppelt, dreifach usw. länger gehört als ein "kürzerer", sondern entsprechend der Funktion $y = ae^{bx}$, die bereits für biologische und physiologische Prozesse bekannt ist.

In Tabelle 2 sind selbstverständlich die empirischen Extremwerte angeführt, die theoretische Funktion ist aus allen aufgetretenen Extensionswerten berechnet (an einem Beispiel werden in Tabelle 4 empirische und theoretische Werte vollständig angeführt).

Tabelle 2: Mittel- und Extremwerte sowie Anpassungsgrad (r) aller Längenangaben (TAKE D, beide Sprecher)

1. Zahl: Sprecherin, 2. Zahl: Sprecher; n' = Zahl der Zeit-/Extensionsklassen, x' basiert auf n'

Symbol	n'	y_{min}	$\bar{x}'$	y_{max}	theoretische Funktion			
					linear ln	$\sqrt[3]{}$	$\sqrt{}$	
i	49	.0298	.1079	.2945	.8986	.9725	.9621	.9512
	40	.0327	.1171	.3073	.9382	.9882	.9824	.9749
t	79	.007	.0826	.307	.9362	.9786	.9884	.9831
	56	.0143	.0808	.3578	.8584	.9819	.9642	.9457
e	79	.02185	.1222	.3055	.9553	.9653	.9785	.9782
	65	.032	.132	.5625	.8012	.9593	.9273	.9026
ε	58	.0256	.1095	.35145	.9036	.9763	.9701	.9596
	37	.0346	.1309	.4993	.8399	.9746	.9496	.9288
ə	79	.0217	.1070	.4352	.8385	.9764	.9515	.9298
	67	.02455	.1091	.2688	.9530	.9925	.9893	.9834
a	65	.02545	.1467	.4352	.9267	.9725	.9777	.9719
	40	.0525	.1489	.4352	.9155	.9724	.9702	.9622
ā	79	.01445	.1156	.362	.9124	.9389	.9580	.9567
	79	.0064	.0971	.3072	.9247	.9265	.9611	.9627
ā	8	.0512	.1399	.22645	.9772	.9507	.9657	.9708
	7	.1024	.1635	.3072	.8573	.9270	.9058	.8943
æ	16	.0403	.08945	.2304	.8728	.9678	.9451	.9300
	3	.0768	.1707	.35605	.9966	.9733	.9831	.9873
o	63	.0064	.09755	.2177	.9679	.9343	.9749	.9810
	43	.022	.1446	.4048	.9356	.9771	.9820	.9766
ɔ	70	.0256	.09949	.2325	.9523	.9904	.9861	.9447
	36	.0346	.1100	.36115	.8461	.9650	.9422	.9237
y	7	.02535	.09703	.2055	.9442	.9565	.9682	.9676
	10	.0609	.1810	.5249	.8567	.9733	.9454	.9266
Y	10	.0298	.0963	.2325	.9430	.9708	.9663	.9624
	4	.0435	.0496	.05625	.9977	.9992	.9988	.9986
u	26	.0256	.0981	.2816	.8945	.9708	.9617	.9504
	16	.0128	.0920	.21825	.9321	.9312	.9550	.9560

(Tabelle 2 Fortsetzung 1)

Symbol	n'	y _{min}	$\bar{x}'$	y _{max}	theoretische Funktion linear ln $\sqrt[3]{\quad}$ $\sqrt{\quad}$			
o	46	.0128	.0959	.2561	.9506	.9877	.9891	.9836
	13	.0165	.0586	.13915	.9186	.9600	.9622	.9568
p	6	.0056	.0471	.11335	.9589	.9674	.9858	.9866
	10	.0107	.0656	.12355	.9841	.9569	.9760	.9815
t	79	.0056	.0674	.3327	.7902	.9649	.9457	.9206
	45	.0145	.0726	.2304	.9100	.9844	.9729	.9616
k	42	.0128	.06825	.23125	.8322	.9393	.9308	.9153
	11	.0261	.06665	.192	.8544	.9651	.9397	.9222
?	79	.0129	.1410	.801	.7904	.9685	.9442	.9152
	54	.0122	.1229	.4252	.9431	.9871	.9949	.9894
b	46	.00275	.06873	.30275	.8686	.9613	.9720	.9562
	33	.00915	.08493	.2815	.9212	.9908	.9897	.9795
d	79	.002	.07027	.7371	.6717	.9667	.9314	.8846
	79	.0032	.06672	.3073	.8361	.9671	.9495	.9271
g	64	.0128	.05317	.21055	.8622	.9852	.9591	.9398
	29	.016	.05485	.1663	.8685	.9689	.9421	.9256
v	79	.0128	.06109	.25545	.8972	.9909	.9778	.9643
	40	.00435	.07955	.5249	.6933	.9238	.9030	.8640
z	35	.0256	.07407	.1792	.9305	.9697	.9677	.9624
	38	.02045	.07842	.3201	.7643	.9348	.8932	.8651
3	-	-	-	-	-	-	-	-
	1	.0426	-	-	-	-	-	-
j	40	.01035	.1133	.4352	.8867	.9934	.9765	.9595
	19	.01615	.0929	.4352	.7520	.9623	.9188	.8839
f	67	.0256	.1298	.768	.6909	.9234	.8753	.8376
	36	.02585	.1373	.5376	.7788	.9309	.9016	.8767
s	79	.0512	.1638	.4603	.8999	.9777	.9630	.9509
	79	.0256	.1201	.4352	.8260	.9466	.9266	.9078
J	50	.0512	.1203	.3548	.8114	.9318	.9049	.8860
	16	.0512	.09725	.1568	.9801	.9846	.9863	.9859
q	74	.02304	.1183	.4224	.8598	.9415	.9409	.9299
	46	.03245	.09119	.1791	.9806	.9845	.9900	.9901
x	32	.02275	.0968	.19845	.9864	.9656	.9803	.9848
	23	.0257	.1049	.4252	.7721	.9462	.9070	.8787

(Tabelle 2 Fortsetzung 2)

Symbol	n'	y _{min}	$\bar{x}'$	y _{max}	theoretische Funktion linear ln $\sqrt[3]{\quad}$ $\sqrt{\quad}$			
h	68	.0128	.3265	2.1632	.8138	.9748	.9769	.9550
	57	.0149	.2102	.789	.9431	.9749	.9909	.9883
l	79	.0088	.0835	.2711	.9542	.9754	.9912	.9897
	48	.0257	.1129	.3201	.9430	.9766	.9822	.9785
R	-	-	-	-	-	-	-	-
	4	.0384	.0521	.06415	.9778	.9661	.9708	.9728
B	44	.00905	.07185	.22535	.9038	.9662	.9694	.9603
	16	.0128	.066	.2098	.9011	.9898	.9775	.9636
D	44	.0506	.14356	.30435	.9717	.9932	.9936	.9908
	58	.019	.14597	.5249	.9063	.9653	.9710	.9637
x	2	.01656	.11255	.06455	1.	1.	1.	1.
	-	-	-	-	-	-	-	-
m	79	.0128	.1452	.801	.7887	.9677	.9407	.9116
	65	.0256	.1899	2.2656	.5529	.9556	.8741	.8042
n	79	.0005	.1277	.7714	.7643	.8103	.9137	.9043
	79	.0135	.1644	1.6869	.5818	.9338	.8732	.8136
q	25	.0385	.1261	.2561	.9611	.9799	.9824	.9799
	13	.0384	.1335	.27995	.9449	.9711	.9729	.9693
/	79	.0081	.8178	1.5345	.2458	.9284	.6717	.5003
	79	.0255	.3148	1.1637	.9168	.9979	.9878	.9750
Pause	78	.02305	3.3916	20.1728	.6969	.9343	.9068	.8616
	76	.1464	4.6775	30.9514	.8324	.9787	.9800	.9583

Die Extensionsanalyse hat einen methodologischen Zweck: Dieser Längenanstieg könnte zwanglos sprechphysiologisch oder akustisch erklärt werden. Der Mehraufwand an Bewegungsaktivität beim Aussprechen eines längeren Lautes vergrößert sich gemäß der biologischen Zunahmefunktion; die zugehörigen akustischen Manifestationen steigen demnach in logarithmischen Zeitabschnitten.

Die akustischen Regelmäßigkeiten wie Vollständigkeit der Schwingungen etc. sind bei der Segmentation so gut wie nicht berücksichtigt worden. Das Segmentierungsprogramm sieht zwar eine Routine zum automatischen Aufsuchen der Nullstellen vor, die der Transkribent zum endgültigen Festlegen der Grenzen benutzen kann; diese wurde für das vorliegende Material sehr selten verwendet.

Weil sich die Quantelungen durch die auditive Festsetzung ergeben haben, ist es einleuchtender, phänomenale Zeitstrukturen und kategoriale Wahrnehmungsprozesse in erster Linie als Erklärung heranzuziehen und erst in zweiter Linie die sprechphysiologischen Abläufe. Von Lauttypen oder Sprechern hängt die Funktion nicht ab; es gibt kaum systematische Unterschiede innerhalb der Lautgruppen bzw. zwischen ihnen (vgl. Tab. 2, rechter Teil).

Verifizierung: Model fitting durch lineare Regression mit semilogarithmischer, semiquadratischer usw. Anpassung. In der Tab. 4 sind als ein Beispiel alle empirischen und theoretischen Werte für [k] und "Standard"-Extensionen angegeben. Die Standardlängen sind aus allen Zeitangaben aller Symbole als häufig wiederkehrende, z.T. mehrfach bei einem Symbol vertretene Extensionen zusammengestellt worden.

Tabelle 3: Anpassung theoretischer Funktionen an die empirische Anstiegsfunktion der Segmentlängen (r)

Funktion	Beispiel [k]	Standardextensionen
$y=ax+b$	.8322	.9497
$y=ae^{bx}$	.9393	.9800
$y=\sqrt[3]{ax^3+bx^2+cx+d}$	.9308	.9927
$y=\sqrt{ax^2+bx+c}$	.9153	.9800
$y=ax^2+bx+c$	.6198	.8550
$y=ax^3+bx^2+cx+d$	.4796	.7845
$y=1+ax+b$	-.8144	-.7502
$y=1/ax^2+bx+c$	-.4100	-.5332

Tabelle 4: Empirische und theoretische Werte der Längenangaben von [k] (TAKE D, Sprecherin) und Standardextensionen

[k]				Standardextensionen			
empirisch	$y=ax+b$	$y=ae^{bx}$	$y=\sqrt[3]{ax^3+bx^2+cx+d}$	empirisch	$y=ax+b$	$y=ae^{bx}$	$y=\sqrt[3]{ax^3+bx^2+cx+d}$
.0128	.0104	.0239	.0208	.0128	.0216	.0250	.0184
.01795	.0132	.0250	.0232	.0256	.0432	.0294	.0243
.02295	.0160	.0261	.0246	.0384	.0648	.0345	.0314
.0256	.0189	.0273	.0260	.0512	.0864	.0404	.0398
.02785	.0217	.0285	.0275	.0548	.1080	.0475	.0495
.03365	.0245	.0298	.0291	.0640	.1279	.0557	.0607
.0345	.0273	.0311	.0308	.0768	.1513	.0654	.0735
.0353	.0302	.0325	.0324	.0896	.1729	.0768	.0879
.0384	.0330	.0339	.0342	.1024	.1945	.0901	.1041
.04015	.0358	.0354	.0360	.1152	.2161	.1058	.1222
.0416	.0386	.0370	.0379	.1280	.2377	.1242	.1423
.0418	.0414	.0387	.0398	.1408	.2593	.1458	.1644
.0425	.0443	.0404	.0418	.1536	.2809	.1712	.1887
.04455	.0471	.0422	.0439	.2045	.3025	.2009	.2154
.0462	.0499	.0441	.0460	.2312	.3241	.2358	.2444
.0548	.0527	.0460	.0482	.2816	.3457	.2768	.2759
.05855	.0556	.0481	.0498	.3328	.3674	.3250	.3100
.05985	.0584	.0502	.0529	.3840	.3890	.3815	.3469
.0635	.0612	.0525	.0553	.4096	.4106	.4478	.3865
.064	.0640	.0548	.0578	.4352	.4322	.5256	.4290
.06625	.0668	.0573	.0603	r = .9497 .9800 .9927			
.06795	.0697	.0598	.0630				
.069	.0725	.0625	.0657				
.07065	.0753	.0653	.0685				
.07105	.0781	.0682	.0714				
.07185	.0810	.0712	.0743				
.07255	.0838	.0744	.0774				
.07275	.0866	.0777	.0805				
.07725	.0894	.0812	.0837				
.0793	.0922	.0848	.0869				
.0794	.0951	.0886	.0903				
.0815	.0979	.0926	.0938				
.08375	.1007	.0967	.0973				
.0846	.1035	.1010	.1009				
.08505	.1063	.1055	.1046				
.0856	.1092	.1102	.1085				
.0896	.1120	.1151	.1123				
.1024	.1148	.1203	.1163				
.10365	.1176	.1256	.1204				
.1112	.1205	.1313	.1246				
.2035	.1233	.1371	.1289				
.23125	.1261	.1432	.1332				
r =				.8322	.9393	.9308	

Die Abweichungen von der ln-Funktion sind teilweise durch niedrige n'-Werte entstanden (z.B. [ø] der Sprecherin, [æ p ʁ] des Sprechers), durch die zu wenig Samples zur Kurvenbestimmung verarbeitet werden können. Teilweise sind die Korrelationskoeffizienten so gering voneinander unterschieden, daß der Unterschied praktisch vernachlässigt werden kann. Das trifft natürlich auch im umgekehrten Fall zu, wenn die ln-Funktion nur geringfügig höher korreliert als andere Kurvenformen. Insgesamt gesehen sind wohl beide Funktionen (ln und 3. Wurzel) gut zur Beschreibung des Längenanstiegs geeignet. Durch die Funktion $y = \sqrt[3]{ax^3+bx^2+cx+d}$ sind im vorliegenden Material die Laute [o h l ɔ ŋ] für beide Sprecher, [u e a a y b ʊ n] der Sprecherin und [o ʔ ʃ] sowie Pausen des Sprechers adäquater bezeichnet. Worin die Ursachen dafür liegen, daß diese oder jene Funktion besser angepaßt ist, d.h., welche phonetischen oder perzeptuellen Gründe zu einem der beiden Kurvenverläufe führen, läßt sich im Material nicht ermitteln.

Die Pausen- und Staupausenlängen beziehen sich sehr viel stärker auf die Zeitverläufe des Gesprächs; sie dokumentieren weniger die apperzeptive Zeitstruktur als vielmehr die Gesprächsdynamik. Das interaktionale Arrangement ist - so können die Zeitrelationen interpretiert werden - demzufolge nicht linear, sondern biologisch-rhythmisch eingeteilt. Die Zeitangaben für Stauungen und Pausen geben die phänomenale Zeitstruktur der Sprecher an, die der Turnorganisation und zeitlichen Gestaltung des Gesprächs zugrundeliegt.

Praktischer Nutzen aus der Extensionsanalyse ist beispielsweise für eine hörgerechte Sprachsynthese oder für das Verstehen von Notaten zu ziehen. Wie die phonetische Notation verwenden auch andere (z.B. konversationsanalytische) Transkriptionssysteme Längenangaben (" : : : " usw.), die sich der Leser des Notats nicht linear, sondern logarithmisch steigend vorzustellen hat. Die Transkription umschreibt nicht "getaktete", sondern rhythmisch aufeinanderfolgende Segmente.

2.2 LAUTHÄUFIGKEIT UND -KOMBINATORIK UMGANGSSPRACHLICH BEEINFLUSSTER AUSSPRACHE

Die Auswertung von Notaten nach phonologischen Kriterien ist methodologisch gerechtfertigt, weil sie nicht im Widerspruch zum Gehalt von Transkriptionen steht. Für phonematische Analysen müssen die Notate nicht unbedingt lautsprachlich oder artikulatorisch orientiert sein. Es gibt dennoch einen Grund, trotzdem auf phonetische Transkripte zurückzugreifen. Mit RICHTER soll zwischen prä- und postrealisatorischen Notaten unterschieden werden; das sind zwei verschiedene Funktionen, mit denen unterschiedliche Verwendungen verbunden sind. ORTMANN (1975) hatte für eine statistische Beschreibung der deutschen Laute und Lautkombinationen das Wörterbuch von KAEDING phonetisch transkribiert. Die Transkripte waren prärealisatorisch und die Resultate sind präskriptiv, nicht deskriptiv aufzufassen. Denn durch die Umwandlung von Schrifttexten in phonetische Symbole ist nicht die lautsprachliche, sondern die schriftkonstituierte Frequenz und Kombinatorik des Deutschen erfaßt, wenngleich die graphematische Komponente weitgehend getilgt ist. Der Status des prärealisatorischen Notats bleibt auch dann erhalten, wenn Schrifttexte so phonetisch transformiert werden, wie die allgemeine Aussprache einer Kombination wäre (z.B. Endsilbe -er als [əʊ]). Ebenso kann eine Analyse der (nicht empirisch entwickelten) Ausspracheempfehlungen des Aussprache-Dudens nicht die in praxi gesprochenen Lautverbindungen beschreiben. Für eine Deskription sind postrealisatorische Notate zu verwenden, die allerdings mit jeglichen phänomenalen und apperzeptiven Vorzügen und Nachteilen belastet sind. Die Analyse befindet sich gewissermaßen auf einem Grenzgebiet, auf dem Kompromisse gemacht werden müssen. Schwer zu entkräften wird ein Argument sein, das aus der wissenssoziologischen Theorie abgeleitet ist: Mit der Konstitution von Transkriptionssystemen ist das gemeinschaftliche Wissen über Lauthäufigkeit und -kombinatorik bereits objektiviert. Jeder Sprachbenutzer weiß - der Fachmann noch auf eine besondere Weise -, daß eine dreifache Konsonantenverbindung nicht "typisch deutsch" ist. Eine statistische Analyse beschreibe nur noch - hyperbolisch formuliert - die Regeln des Notationssystems und ob der Transkribent die Regeln einhält. Der Transkribent weiß vorher, daß sich

das Rundungszeichen [°] nur mit Vokalen kombinieren kann; er verwendet auch keine anderen Kombinationen und die statistische Analyse bestätigt, daß es keine "gerundeten Konsonanten" gibt. Die Notwendigkeit und der Einfluß des Vorwissens werden in der Phonetik nicht nur akzeptiert, sondern gewöhnlich auch gefordert: Der Transkribent muß die Sprache beherrschen, d.h. Lauthäufigkeit und -kombinatorik per erworbener Kompetenz kennen, und gelernt haben, was mit einem Symbol gemeinschaftlich konstituiert und objektiviert worden ist, wie es zu "klingen" hat usw., bevor er phonetisch transkribieren kann. Eine Frequenz- und Kombinationsanalyse ist in diesem Sinne keine "Vorstufe für experimentelle Untersuchungen des Sprechbewegungsablaufes" (LINDNER 1975, 77), sondern eine Nachbereitung des interindividuellen Wissens über statistische Regelmäßigkeiten von Artikulationsbewegungen, die der gesprochenen Sprache zugrundeliegen.

Der Gewinn bei der Analyse postrealisatorischer Notate liegt vor allem in jenen Abweichungen, die sich beim "Auflegen" der intersubjektiv verbindlichen Wissensfolie über idealtypische Lautverbindungen auf konkrete, empirische und teilweise vortypische Varianten ergeben. Die statistische Analyse regt zur Überprüfung an, ob der allgemeine Wissensvorrat auch spezielle Unterarten und Besonderheiten umfaßt, ob er erweitert werden muß, wie hoch das Idealisierungsniveau ist und welches Resultat entsteht, wenn eine objektive Wissensstruktur mit empirischen Befunden in Verbindung gebracht wird (beispielsweise durch Transkribieren auf das empirische Sprachmaterial "aufgeprägt" wird).

Am günstigsten ist demnach postrealisatorisches Transkriptmaterial, das mit einem gemeinschaftlich akzeptierten und prototypischen Notationssystem angefertigt worden ist und bei dem der Transkribent zwar korrekt, aber "unorthodox" verfahren ist (also auch "ungewöhnliche" Laute und Kombinationen zugelassen hat). Das phonematische Vorwissen kann zwar nicht ausgeschaltet werden (soll es auch nicht), es kann aber in gewisser Weise relativiert werden, indem phonetische Transkripte für phonologische Analysen verwendet werden. Die phonetische Notation nimmt weniger Rücksicht auf systematische Kombinationen. Phonetisch werden z.B. auch Folgen mehrerer Glottiseinsätze, mehrerer Konsonanten oder VVV-Ketten (statt der standardgemäßen VCV-/CVC-Verbindungen) notiert, wie sie in realen Sprachsequenzen auftreten. Die phonetische Intention des Tran-

s-kribierens (neben phonematischen auch lautsprachliche und sprechphysiologische Einheiten zu dokumentieren) wird zum methodischen Vorteil für die phonologische Analyse. Der Transkribent muß zu diesem Zweck nicht unabdingbar native speaker sein; er muß nur das Notationssystem anwenden können (in wissenssoziologischer Formulierung: er muß Teilhaber des gesellschaftlichen Wissens sein). Das statistische Resultat - das zeigt den Kompromiß - ist sowohl phonetisch als auch phonologisch und, im größeren methodologischen Zusammenhang, auch wissenstheoretisch interpretierbar; wie jede Statistik keine bestimmte inhaltliche Auslegung zwingend vorschreibt.

Die Aufnahmen des TAKE D wurden phonetisch mit dem IPA-System¹ computerunterstützt transkribiert. Der Transkribent beherrschte weder die berliner noch die schwäbisch-alemannische Umgangssprache. Alle Lautvarianten, die mit der speziellen Technik des auditiv-visuellen Segmentierens und Transkribierens zu identifizieren waren, wurden transkribiert, auch wenn sie ungewohnt anmuteten. Die Aufnahmen wurden in zwei Versionen transkribiert: auditiv oder akustisch (d.i. visuell anhand von Kurvenformen) orientiert. Beide Versionen werden zur Analyse herangezogen: Die Frequenzanalyse stützt sich auf die auditiv, die Distributionsanalyse auf die akustisch orientierte Transkription.

Grundlage sind alle verwendeten Transkriptionssymbole, Sonderzeichen mit rein phonetischer Bedeutung inbegriffen. Die statistische Analyse verwendet die Verfahren, die ALTMANN/LEHFELDT (1980) entwickelt haben und die eine umfassende Beschreibung und Klassifikation von Lautfrequenzen und -kombinationen ermöglichen.

Das Prozedere statistischer Auswertungen ist normalerweise sehr aufwendig. Die Häufigkeiten und Kombinationen müssen manuell ausgezählt und verrechnet werden; bei Verwendung einer Rechenanlage (das besagt üblicherweise der Begriff "computerunterstützt") sind alle Auszählungen einzulesen, wie z.B. bei dem Verfahren von WOTHKE (1980). Das vom Münchener Institut entwickelte Transkriptionsprogramm enthebt den Analytiker aller zeitraubenden Dateneingaben; während des Segmentierens wird automatisch ein Transkriptions-File angelegt, der problemlos ausgewertet werden kann. Solche Programme sind jedoch noch nicht sehr verbreitet, weil vor allem die kostspielige Hardware fehlt. Um zu demonstrieren, wie eine phonologische Statistik mit einfacher Hardware durchgeführt werden kann, wurden die nachfolgenden Analysen mit einem programmierbaren Taschenrechner (TI 59 mit PC 100C) durchgerechnet. Die zugehörige Software, die viele Parameter aus der "Quantitativen Phonologie" von ALTMANN/LEHFELDT (1980) berechnet, ist im Anhang aufgelistet.

Die Transkriptions-Files des TAKE D wurden zu diesem Zweck teilweise maschinell vorausgewertet (z.B. automatische Umsetzung des computerinternen phonetischen Codes in normale phonetische Symbole), teilweise manuell ausgezählt (z.B. nach dem Schema S. 245 in ALTMANN/LEHFELDT 1980).²

2.2.1 Frequenzanalyse

In Tabelle 5 sind die statistischen Parameter der Häufigkeiten der umgangssprachlich beeinflussten Aussprache aufgeführt. Die verwendete Stichprobe hatte "natürliche" Grenzen und konnte nicht beliebig erweitert werden. Schätzt man die benötigte Größe der Stichprobe aus der geringsten Häufigkeit (z.B. $[3]_{\text{Sprecher}} = .00035$), dann müßte $N = 1085252.128$ betragen. Nach der Faustregel: $1000 \times (2 \times \text{Inventar})$ wären 1000000 Messungen nötig.

Die Parameter bedeuten (vgl. ALTMANN/LEHFELDT 1980, 142-182):

	verwendete Formel
K	Symbolinventar
N	Gesamthäufigkeiten
$\hat{p}_i$	relative Häufigkeit f_i/N
R	Repeatrate $(\chi^2 + N)/NK$
D	relativiertes, R bzw. χ^2 $\chi^2/N(K-1)$
V	Variationsquotient $\sqrt{\chi^2/N}$
H_0, H_1	Entropie $\text{ld}K; \approx H_0 - (\chi^2/(2N \ln 2))$
h	relative Entropie H_1/H_0
Re	relative Redundanz $1-h$
I_L	Index Funktionalmittelwerte 2^{-H_1}

Im Vergleich des Parameters $\hat{R}$ (geschätzter R-Wert) mit den Ergebnissen anderer Analysen der deutschen Sprache ergeben sich Differenzen, die durch die verschiedenen Beträge von K entstehen, worin sich der Unterschied von phonetischen und phonematischen Einheiten zeigt (die ersten drei Werte sind der Zusammenstellung von ALTMANN/LEHFELDT 1980, 154f. entnommen):

Material	K	$\hat{R}$	Autor
deutscher Standard	29	.072007	Ovčinnikov 1962
deutscher Standard	30	.073938	Fitiaľova 1965
deutscher Standard	33	.059241	Kučera, Monroe 1968
umgangssprachlich beeinflusst	50	.035951	
berlinisch beeinflusst	48	.025182	
schwäbisch-alemannisch beeinfl.	47	.037431	

Tabelle 5: Parameter und Häufigkeiten der umgangssprachlich beeinflussten Aussprache

Symbol	berliner und schwäbisch-alemannischer Einfluß f_i	$\hat{p}_i$	schwäbisch-alemannisch f_i	$\hat{p}_i$	berliner Einfluß f_i	$\hat{p}_i$
3	1	.0001298533	-	-	1	.000354
X	2	.0002597065	2	.000410	-	-
+	2	.0002597065	-	-	2	.000710
R	4	.0005194131	-	-	4	.001416
~	5	.0006492663	3	.000615	2	.000710
o	12	.0015582392	12	.002461	-	-
Y	14	.0018179457	10	.002051	4	.001416
ø	15	.0019477990	8	.001641	7	.002478
p	16	.0020776523	6	.001231	10	.003540
Y	19	.0024672121	7	.001436	12	.004248
æ	22	.0028567718	19	.003897	3(-1)	.001062
η	47	.0061031035	34	.006973	13	.004602
u	53	.0068822231	33	.006768	20	.007079
k	59	.0076613427	47	.009639	12	.004248
x	73	.0094792884	43	.008819	30	.010619
j	75	.0097389949	52	.010664	23	.008142
ſ	81	.0105181145	61	.012510	20	.007079
o	86	.0111673809	69	.014151	17	.006018
ø	88	.0114270874	63	.012920	25	.008849
b	105	.0136345929	55	.011279	50	.017699
z	105	.0136345929	48	.009844	57	.020177
h	109	.014154006	88	.018048	21	.007434
g	115	.0149331256	78	.015997	37	.013097
l	121	.0157122452	63	.012920	58	.020531
ε	122	.0158420984	77	.015792	45	.015929
.	127	.0164913648	75	.015381	52	.018407
a	129	.0167510713	80	.016407	49	.017345
v	132	.0171406311	60	.012305	72	.025487
o	138	.0179197507	88	.018048	50	.017699
o	149	.0193481366	86	.017637	63	.022301
f	150	.0194779899	100	.020509	50	.017699
ll	158	.0205168160	82	.016817	76	.026903

Tabelle 5 (Fortsetzung)

Symbol	berliner und schwäbisch-alemannischer Einfluß		schwäbisch-alemannisch		berliner Einfluß	
	f _i	p _i	f _i	p _i	f _i	p _i
h	160	.0207765225	92	.0188679	68	.024071
v	178	.0231138812	122	.0250205	56	.019823
l	192	.0249318270	105	.0215340	87	.030796
ç	193	.0250616803	128	.0262510	65	.023009
t	215	.0279184521	151	.0309680	64	.022655
ʔ	222	.0288274250	159	.0326087	63	.022301
e	225	.0292169848	133	.0272765	92	.032567
•	231	.0299961044	154	.0315833	77	.027257
ː	234	.0303856642	100	.0205086	134	.047434
/	235	.0305155175	126	.0258409	109	.038584
l	251	.0325931697	164	.0336341	87	.030796
m	253	.0328528762	153	.0313782	100	.035389
ə	367	.0476561486	258	.0529122	109	.038584
d	411	.0533696922	253	.0518868	158	.055929
s	483	.0672191274	324	.0664479	159	.056283
a	488	.0633683937	314	.0643970	174	.061593
n	489	.0634982470	315	.0646021	174	.061593
• dev.	540	.0701207635	376	.0771124	164	.058053
K	50		47		48	
N	7701		4876		2825	
ɛf ²	2132101		889940		280773	
χ ²	6142.013894		3702.174733		1945.656283	
R	.0359512113		.0374311641		.035181831	
D	.0162767462		.0165057546		.0146537848	
V	.893062463		.8713579698		.8298963115	
H ₀	5.64385619		5.554588852		5.584962501	
H ₁	5.068537855		5.006895135		5.088150347	
h	.8980628997		.9013979735		.911044675	
Re	.1019371003		.0986020265		.088955325	
I _L	.029800124		.0311010019		.0293977521	

NB: Diphthonge und Affrikata sind nach ihren Bestandteilen gezählt
dev. = devoiced

Merkmale der Umgangssprache sind häufig auftretende Entstimmlichungen (540; aus 376 und 164) und Zusammenziehungen von Silben, wobei der Restlaut (zumeist Nasal) silbisch gesprochen wird (mittlere Häufigkeit, 75). Schwäbisch-alemannisch sind die Vokalrundungen (12), die nicht im berliner Einfluß auftreten. Die berlinisch beeinflusste Aussprache zeigt sich u.a. in der höheren Zahl von vokalischen Ersatzbildungen des R-Lautes (72 gegenüber 60).

2.2.2 Distributionsanalyse

Mit den statistischen Maßen der Distribution werden die Kombinationen von Transkriptionssymbolen ohne Rücksicht auf die Frequenz beschrieben. Es wird die Menge von Symbolen ausgewertet, die als Vorgänger oder Nachfolger eines Symbols i auftritt. Es hängt vom methodologischen Standpunkt ab, ob diese Kombinatorik als phonetisch-artikulatorische oder phonematische oder aber als rein transkriptionelle Relation interpretiert wird. Die statistischen Begriffe wie Attraktivität oder Aggressivität besagen lediglich, daß sich bestimmte Symbole mit größeren oder kleineren Mengen anderer Symbole kombinieren. In einem begrenzten Material lassen sich nicht alle statistisch möglichen Kombinationen finden, auch nicht alle Verbindungen, die von einem Transkriptionssystem vorgesehen sind. Die Maßzahlen geben in jeglicher Hinsicht nur Ausschnitte an; selektiv aufgetretene Beziehungen innerhalb von Transkriptionsregeln, innerhalb des Phonemsystems oder lautsprachlichen Realisierung. Die Schätzungen und Klassifikationen aufgrund des vorliegenden Materials können nur Annäherungen sein. Ein Problem der statistischen Auswertungen, das auch von ALTMANN/LEHFELDT (1980, 217-220) diskutiert wird, entfällt durch die Vorentscheidung für ein bestimmtes Transkriptionssystem: die Wahl von Rahmeneinheiten wie Phonem, phonotaktische Einheit, Silbe etc.; mit den phonetischen Symboleinheiten ist die Ebene der Segmentierung vorgegeben. Auch aus diesem Grund ist es angebracht, die statistischen Resultate zunächst nur als Ausdruck vom Symbol-Relationen aufzufassen.

In den Tabellen 6.1 - 6.3 sind die Distributionsmaße für die Symbolfolgen angeführt, die bei der Transkription der Umgangssprache verwendet worden sind. Von den Parametern und Algorithmen, die

ALTMANN/LEHFELDT (1980, 228ff.) vorschlagen, sind diejenigen ausgewählt worden, die für Einzelfälle am ehesten geeignet sind (Totalmaße für die Beschreibung einer Sprache insgesamt sind nicht verwendet worden) bzw. die den Aufwand an Datenarrangement minimal halten. Die verwendete Software ist so eingerichtet, daß aus drei Werten und K alle anderen Parameter berechnet werden können. Bei alternativen Rechenwegen bzw. Algorithmen wurden die für das Rechenprogramm ökonomischeren benutzt (siehe folgende Aufstellung):

Bezeichnung	Algorithmus
A_i	Menge der Vorgänger
iB	Menge der Nachfolger
VE	$= T_i$; Vereinigungsmenge von A_i und iB
DU	$A \cap B$
SD	$A \oplus B$
Maße der Assoziativität	
At	Attraktivität $At(i)=[A_i]/K$
Ag	Aggressivität $Ag(i)=[iB]/K$
As	Assoziativität $As(i)=[A_i \cup iB]/K$
As^*	Maß der Paarbildung von i $As^* = ([A_i] + [iB] - G(i)) / (2K - 1)$
At_i	interne Attraktivität $At_i = [A_i] / VE$
Ag_i	interne Aggressivität $Ag_i = [iB] / VE$
Maße der Symmetrie	
S	Symmetrie $S = [A_i \cap iB] / K$
S_i	interne Symmetrie $S_i = [A_i \cap iB] / VE$
S_i^*	potentielle Symmetrie $S_i^* = 2[A_i \cap iB] / ([A_i] + [iB])$
S'	Antisymmetrie $S' = [A_i \oplus iB] / K$
S'_i	interne Antisymmetrie $S'_i = [A_i \oplus iB] / VE$
$S_i'^*$	potentielle Antisymmetrie $S_i'^* = [A_i \oplus iB] / ([A_i] + [iB])$
Reflexivität:	
G	Kombination i mit i $= 1$, wenn ii auftritt, sonst $= 0$

Die Symbolfolgebeziehungen wurden nach dem Kriterium für $K = 50$ klassifiziert; Werte zwischen 18 und 32 gelten als semi-assoziativ (siehe Tabellen 7.1 bis 7.3).

Tabelle 6.1: Distributionsmaße im Längen-Potenzialmodell

Sy	Ai	iB	VE	DU	SD	At	Ag	AS	As*	At _i	Ag _i	S	S _i	S _i *	S'	S' _i	S _i '*	G
3	1	1	2	0	2	.02	.02	.04	.02020	.05	.05	0	0	0	.04	1	1	0
χ	3	2	5	0	5	.06	.04	.01	.05051	.06	.04	0	0	0	.01	1	1	0
κ	2	2	2	0	0	.04	.04	.04	.04040	1	1	.04	1	1	0	0	0	0
κ	6	5	10	1	9	.12	.01	.02	.11111	.06	.05	.02	.01	.181818	.18	.9	.818182	0
~	4	4	8	0	8	.08	.08	.16	.08081	.05	.05	0	0	0	.16	1	1	0
o	4	10	12	2	10	.08	.2	.24	.14141	.33333	.83333	.04	.166667	.285714	.2	.833333	.714286	0
Y	11	10	16	5	11	.22	.2	.32	.21212	.6875	.625	.1	.3125	.476190	.22	.6875	.523809	0
φ	8	8	15	1	14	.16	.16	.3	.16162	.53333	.53333	.02	.066667	.125	.28	.933333	.875	0
p	11	10	18	3	15	.22	.2	.36	.21212	.61111	.555556	.06	.166667	.285714	.3	.833333	.714286	0
Y	15	17	25	7	18	.3	.34	.5	.31313	.6	.68	.14	.28	.4375	.36	.72	.5625	1
π	13	16	24	5	19	.26	.32	.48	.29293	.541667	.666667	.1	.208333	.344828	.38	.791667	.655172	0
q	19	16	27	8	19	.38	.32	.54	.34343	.703704	.592593	.16	.296296	.457143	.38	.703704	.542857	1
u	26	13	29	10	19	.52	.26	.58	.39394	.896551	.448276	.2	.344828	.512821	.38	.655172	.487179	0
k	22	23	35	10	25	.44	.46	.7	.45455	.628571	.657143	.2	.285714	.444444	.5	.714286	.555556	0
x	32	16	34	14	20	.64	.32	.68	.484848	.941176	.470588	.28	.411765	.583333	.4	.588235	.416667	0
j	22	25	30	17	13	.44	.5	.6	.464646	.733333	.833333	.34	.566667	.723404	.26	.433333	.276596	1
f	30	20	37	13	24	.6	.4	.74	.505051	.810810	.540541	.26	.351351	.52	.48	.648649	.48	0
q	14	24	29	9	20	.28	.48	.58	.383838	.482759	.827586	.18	.310345	.473684	.4	.689655	.526316	0
a	22	27	35	14	21	.44	.54	.7	.494949	.628571	.771429	.28	.4	.571429	.42	.6	.428571	0
b	30	27	38	19	19	.6	.54	.76	.575758	.789474	.710526	.38	.5	.666667	.38	.5	.333333	0
z	21	23	33	11	22	.42	.46	.66	.444444	.636364	.696969	.22	.333333	.5	.44	.666667	.5	0
h	29	17	37	9	28	.58	.34	.74	.464646	.783784	.459459	.18	.243243	.391304	.56	.756757	.608696	0
g	28	30	39	19	20	.56	.6	.78	.585859	.717949	.769231	.38	.487179	.655172	.4	.512821	.344828	0
i	35	21	38	18	20	.7	.42	.76	.555556	.921053	.552632	.36	.473684	.642857	.4	.526316	.357143	1

Tabelle 6.1 (Fortsetzung)

Sy	Al	iB	VE	DU	SD	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S' _i	S' _i *	G
e	30	32	37	25	12	.6	.64	.74	.616162	.810811	.864865	.5	.675676	.806452	.24	.324324	.193548	1
f	7	32	33	6	27	.14	.64	.66	.393939	.212121	.969697	.12	.181818	.307692	.54	.818182	.692308	0
g	36	27	42	21	21	.72	.54	.84	.636364	.857143	.642857	.42	.5	.666667	.42	.5	.333333	0
h	32	21	38	15	23	.64	.42	.76	.535354	.842105	.552632	.3	.394737	.566038	.46	.605263	.433962	0
i	27	27	38	16	22	.54	.54	.76	.535354	.710526	.710526	.32	.421053	.592593	.44	.578947	.407407	1
j	37	31	40	28	12	.74	.62	.8	.676768	.925	.775	.56	.7	.823529	.24	.3	.176471	1
k	35	33	41	27	14	.7	.66	.82	.676768	.853659	.804878	.54	.658537	.794118	.28	.341463	.205882	1
l	32	26	39	19	20	.64	.52	.78	.585859	.820513	.666667	.38	.487179	.655172	.4	.512821	.344828	0
m	35	33	42	26	16	.7	.66	.84	.676768	.833333	.785714	.52	.619048	.764706	.32	.380952	.235294	1
n	27	31	39	19	20	.54	.62	.78	.575758	.692308	.794872	.38	.487179	.655172	.4	.512821	.344828	1
o	33	29	37	25	12	.66	.58	.74	.616162	.891892	.783784	.5	.675676	.806452	.24	.324324	.193548	1
p	37	32	43	26	17	.74	.64	.86	.686869	.860465	.744186	.52	.604651	.753623	.34	.395349	.246377	1
q	37	31	42	26	16	.74	.62	.84	.676768	.880952	.738095	.52	.619048	.764706	.32	.380952	.235294	1
r	25	37	44	18	26	.5	.74	.88	.616162	.568182	.840909	.36	.409091	.580645	.52	.590909	.419355	1
s	36	32	42	26	16	.72	.64	.86	.676768	.857143	.761905	.52	.619048	.764706	.32	.380952	.235294	1
t	37	22	43	16	27	.74	.44	.86	.595960	.860465	.511628	.32	.372093	.542373	.54	.627907	.457627	0
u	41	36	43	34	9	.82	.72	.86	.767677	.953488	.837209	.68	.790698	.883117	.18	.209302	.116883	1
v	40	30	43	27	16	.8	.6	.86	.696970	.930233	.697674	.54	.627907	.771429	.32	.372093	.228571	1
w	26	35	38	23	15	.52	.7	.76	.616162	.684211	.921053	.46	.605263	.754098	.3	.394737	.245902	0
x	38	37	43	32	11	.76	.74	.86	.747475	.883721	.860465	.64	.744186	.853333	.22	.255814	.146667	1
y	39	40	47	32	15	.78	.8	.94	.787879	.829787	.851064	.64	.680851	.810127	.3	.319149	.189873	1
z	39	40	45	34	11	.78	.8	.9	.787879	.866667	.888889	.68	.755556	.860759	.22	.244444	.139241	1
aa	40	34	43	31	12	.8	.68	.86	.737374	.930233	.790698	.62	.720930	.837838	.24	.279069	.162163	1
ab	37	37	42	32	10	.74	.74	.84	.747475	.880952	.880952	.64	.761905	.864865	.2	.238095	.135135	0
ac	43	42	46	39	7	.86	.84	.92	.848485	.934783	.913043	.78	.847826	.917647	.14	.152174	.082353	1
ad	14	40	40	14	26	.28	.8	.8	.545455	.45	.1	.28	.45	.518519	.52	.65	.481481	0

Sy	Al	iB	VE	DU	SD	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S' _i	S' _i *	G
ae	3	2	5	0	5	.063830	.042553	.106383	.053763	.6	.4	0	0	0	.106383	1	1	0
af	4	3	7	0	7	.085106	.063829	.148936	.075269	.571429	.428571	0	0	0	.148936	1	1	0
ag	8	7	14	1	13	.170213	.148936	.297872	.161290	.571429	.5	.021277	.071429	.133333	.276596	.928571	.866667	0
ah	8	8	11	5	6	.170213	.170213	.234043	.161290	.727273	.727273	.106382	.454545	.625	.127659	.545455	.375	1
ai	7	6	12	1	11	.148936	.127659	.255319	.139785	.583333	.5	.021277	.083333	.153846	.234043	.916667	.846154	0
aj	8	6	12	2	10	.170213	.127659	.255319	.150537	.666667	.5	.042553	.166667	.285714	.212766	.833333	.714286	0
ak	4	10	13	1	12	.085106	.212766	.276596	.150538	.307692	.769231	.021277	.076923	.142857	.255319	.923077	.857142	0
al	13	15	23	5	18	.276586	.319149	.489362	.301075	.565217	.652174	.106383	.217391	.357142	.382979	.782609	.642857	0
am	21	11	26	6	20	.446809	.234043	.553191	.344086	.807692	.423077	.127659	.230769	.375	.425532	.769231	.625	0
an	18	14	25	7	18	.382979	.297872	.531915	.344086	.72	.56	.148936	.28	.4375	.382979	.72	.5625	0
ao	26	12	30	8	22	.553191	.255319	.638298	.408602	.866667	.4	.170213	.266667	.421053	.468085	.733333	.578947	0
ap	24	19	35	8	27	.510638	.404255	.744681	.462366	.685714	.542857	.170213	.228571	.372093	.574468	.771429	.627907	0
aq	16	15	25	6	19	.340426	.319149	.531915	.333333	.64	.6	.127659	.24	.387097	.404255	.76	.612903	0
ar	19	19	27	11	16	.404255	.404255	.574468	.408602	.703704	.703704	.234043	.407407	.578947	.340426	.592593	.421053	0
as	24	22	31	15	16	.510638	.468085	.659574	.494624	.774194	.709677	.319149	.483871	.652174	.340426	.516129	.347826	0
at	25	17	36	6	30	.531915	.361702	.765957	.451613	.694444	.472222	.127659	.166667	.285714	.638298	.833333	.714286	0
au	25	16	32	9	23	.531915	.340426	.680851	.440860	.78125	.5	.191489	.28125	.439024	.489362	.71875	.560976	0
av	31	18	34	15	19	.659574	.382979	.723404	.526882	.911765	.529412	.319149	.441176	.612245	.404255	.558824	.387755	0
aw	20	25	33	12	21	.425532	.531915	.702128	.483871	.606061	.757576	.255319	.363637	.533333	.446809	.636364	.466667	0
ax	14	23	28	9	19	.297872	.489362	.595745	.397849	.5	.821485	.191489	.321429	.486487	.404255	.678571	.513512	0
ay	13	32	34	11	23	.276596	.680851	.723404	.483871	.382353	.941176	.234043	.323529	.488889	.489362	.676471	.511111	0
az	23	23	31	15	16	.489362	.489362	.659574	.494624	.741935	.741935	.319149	.483871	.652174	.340426	.516129	.347826	0
ba	28	24	34	18	16	.595745	.510638	.723404	.559139	.823529	.705882	.382979	.529412	.692308	.340426	.470589	.307692	0
bb	31	22	38	15	23	.659574	.468085	.808511	.569892	.815789	.578947	.319149	.394737	.566038	.489362	.605263	.433962	0

Tabelle 6.2 (Fortsetzung)

Sy	Al	iB	VE	DU	SD	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S _i '	S _i *	G
11	26	21	33	14	19	.553191	.446809	.702128	.505376	.787879	.636364	.297872	.424242	.595744	.404255	.575758	.404255	o
o	32	28	39	21	18	.680851	.595744	.829787	.634409	.820513	.717949	.446809	.538462	.7	.382979	.461538	.3	1
h	28	13	38	3	35	.595745	.276596	.808511	.440860	.736842	.342105	.063829	.078947	.146341	.744681	.921053	.853659	o
o	26	22	35	13	22	.553191	.468085	.744681	.505376	.742857	.628571	.276595	.371429	.541667	.468085	.628571	.458333	1
h	29	26	36	19	17	.617021	.553191	.765957	.580645	.805556	.722222	.404255	.527778	.690909	.361702	.472222	.360909	1
f	32	33	41	24	17	.680851	.702128	.872340	.688172	.780488	.804878	.510638	.585366	.738462	.361702	.414634	.261538	1
:	36	28	39	25	14	.765957	.595745	.829787	.688172	.923077	.717949	.531915	.641026	.78125	.297872	.358974	.21875	o
l	27	23	31	19	12	.574468	.489617	.659574	.537634	.870968	.741935	.404255	.612903	.76	.255319	.387097	.24	o
v	27	30	39	18	21	.574468	.638298	.829787	.602151	.692308	.769231	.382979	.461538	.631579	.446809	.538462	.368421	1
/	30	27	37	20	17	.638298	.574468	.787234	.602151	.810811	.729730	.425532	.540541	.701754	.361702	.459460	.298246	1
q	33	29	40	22	18	.702128	.617021	.851064	.655914	.825	.725	.468085	.55	.709677	.382979	.45	.29032	1
e	30	22	36	16	20	.638298	.468085	.765957	.548387	.833333	.611111	.340426	.444444	.615385	.425532	.555556	.384615	1
t	36	30	42	24	18	.765957	.638298	.893617	.698925	.857143	.714286	.510638	.571429	.727273	.382979	.428571	.272727	1
m	36	37	45	28	17	.765957	.787234	.957447	.774194	.8	.822222	.595745	.622222	.767123	.361702	.377778	.232877	1
.	33	20	40	13	27	.702128	.425532	.851064	.569892	.825	.5	.276596	.325	.490566	.574468	.675	.509434	o
y	26	32	40	18	22	.553191	.680851	.851064	.612903	.65	.8	.382979	.45	.620689	.468085	.55	.379310	1
l	25	27	33	19	14	.531915	.574468	.702128	.559139	.757576	.818182	.404255	.575758	.730769	.297872	.424242	.269308	o
d	35	36	40	31	9	.744681	.765957	.851064	.752688	.875	.9	.659574	.775	.873239	.191489	.225	.126761	1
a	36	37	44	29	15	.765957	.787234	.936170	.774194	.818182	.840909	.617021	.659091	.794521	.319149	.340909	.205479	1
a	32	37	40	29	11	.680851	.787234	.851064	.741935	.8	.925	.617021	.725	.840579	.234043	.275	.159420	o
n	40	39	44	35	9	.851064	.829787	.936170	.838710	.909091	.886364	.744681	.795455	.886076	.191489	.204545	.113924	1
s	39	33	44	28	16	.829787	.702128	.936170	.763441	.886364	.75	.595745	.636364	.777778	.340425	.363636	.222222	1
o	11	37	39	9	30	.234043	.787234	.829787	.516129	.282051	.948718	.191489	.230769	.375	.638298	.769231	.625	o

Sy	Al	iB	VE	DU	SD	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S _i '	S _i *	G
3	1	1	2	0	2	.020833	.020833	.041667	.021053	.5	.5	o	o	o	.041667	1	1	o
~	1	1	2	0	2	.020833	.020833	.041667	.021053	.5	.5	o	o	o	.041667	1	1	o
+	2	2	2	2	0	.041667	.041667	.041667	.042105	1	1	.041667	1	1	o	o	o	o
e	3	3	6	0	6	.0625	.0625	.125	.063158	.5	.5	o	o	o	.125	1	1	o
r	5	5	9	1	8	.104167	.104167	.1875	.105263	.555556	.555556	.020833	.111111	.2	.166667	.888889	.8	o
y	5	4	8	1	7	.104167	.083333	.166667	.094737	.625	.5	.020833	.125	.222222	.145833	.875	.777778	o
ø	5	6	10	1	9	.104167	.125	.208333	.115789	.5	.6	.020833	.1	.181818	.1875	.9	.818182	o
p	6	4	9	1	8	.125	.083333	.1875	.105263	.666667	.444444	.020833	.111111	.2	.166667	.888889	.8	o
k	13	10	19	4	15	.270833	.208333	.395833	.242105	.684211	.526316	.083333	.210526	.347826	.3125	.789474	.652174	o
y	12	9	19	2	17	.25	.1875	.395833	.221053	.631579	.473684	.041667	.105263	.190476	.354167	.894737	.809524	o
q	7	12	17	2	15	.145833	.25	.354167	.189474	.411765	.705882	.041667	.117647	.210526	.3125	.882353	.789474	1
o	9	7	12	4	8	.1875	.145833	.25	.168421	.75	.583333	.083333	.333333	.5	.166667	.666667	.5	o
j	11	18	24	5	19	.229167	.375	.5	.305263	.458333	.75	.104167	.208333	.344828	.395833	.791667	.655172	o
u	10	13	20	3	17	.208333	.270833	.416667	.242105	.5	.65	.0625	.15	.260869	.354167	.85	.739130	o
h	14	14	24	4	20	.291667	.291667	.5	.294737	.583333	.583333	.083333	.166667	.285714	.416667	.833333	.714286	o
j	19	13	25	7	18	.395833	.270833	.520833	.326316	.76	.52	.145833	.28	.4375	.375	.72	.5625	1
ø	12	11	20	3	17	.25	.229167	.416667	.242105	.6	.55	.0625	.15	.260869	.354167	.85	.739130	o
x	22	10	25	7	18	.458333	.208333	.520833	.336842	.88	.4	.145833	.28	.4375	.375	.72	.5625	o
g	22	16	27	11	16	.458333	.333333	.5625	.4	.814815	.592593	.229167	.407407	.578947	.333333	.592593	.421053	o
e	24	23	31	16	15	.5	.479167	.645833	.484211	.774194	.741935	.333333	.516129	.680851	.3125	.483871	.319149	1
a	19	25	32	12	20	.395833	.520833	.666667	.463158	.59375	.78125	.25	.375	.545455	.416667	.625	.454545	o
b	21	21	30	12	18	.4375	.4375	.625	.442105	.7	.7	.25	.4	.571429	.375	.6	.428571	o
f	22	27	32	17	15	.458333	.5625	.666667	.505263	.6875	.84375	.354167	.53125	.693878	.3125	.46875	.306122	1
o	19	20	29	10	19	.391833	.416667	.604167	.410526	.655172	.689655	.208133	.344828	.512821	.195833	.655172	.481179	o

Tabelle 6.3 (Fortsetzung)

SY	AI	iB	VE	DU	SD	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S' _i	S' _i *	G
•	23	4	26	1	25	.479167	.083333	.541667	.284211	.884615	.153846	.020833	.384615	.074074	.520833	.961538	.925926	o
v	21	13	25	9	16	.4375	.270833	.520833	.357895	.84	.52	.1875	.36	.529412	.333333	.64	.470588	o
z	21	17	26	12	14	.4375	.354167	.541667	.4	.807692	.653846	.25	.461538	.631579	.291667	.538462	.368421	o
i	13	29	32	10	22	.270833	.604167	.666667	.431579	.40625	.90625	.208333	.3125	.476190	.458333	.6875	.523809	1
?	20	19	26	13	13	.416667	.395833	.541667	.4	.769231	.730769	.270833	.5	.666667	.270833	.5	.333333	1
o	18	29	36	11	25	.375	.604167	.75	.494737	.5	.805556	.229167	.305556	.468085	.520833	.694444	.531915	o
t	23	26	33	16	17	.479167	.541667	.6875	.515789	.696969	.787879	.333333	.484849	.653061	.354167	.515152	.346939	o
q	19	27	30	16	14	.395833	.5625	.625	.473684	.633333	.9	.333333	.533333	.695652	.291667	.466667	.304348	1
h	25	24	32	17	15	.520833	.5	.666667	.505263	.78125	.75	.354167	.53125	.693878	.3125	.46875	.306122	1
p	18	30	37	11	26	.375	.625	.770833	.505263	.486487	.810811	.229167	.297297	.458333	.541667	.702703	.541667	o
l	22	25	34	13	21	.45833	.520833	.708333	.494737	.647059	.735294	.270833	.382353	.553191	.4375	.617647	.446809	o
.	13	28	31	10	21	.270833	.583333	.645833	.431579	.419355	.903226	.208333	.322581	.487805	.4375	.677419	.512195	o
l	27	19	34	12	22	.5625	.395833	.708333	.484211	.794118	.558824	.25	.352941	.521739	.458333	.647059	.478261	o
l	24	26	33	17	16	.5	.541667	.6875	.515789	.727273	.787879	.354167	.515152	.68	.333333	.484848	.32	1
e	24	27	33	18	15	.5	.5625	.6875	.526316	.727273	.818182	.375	.545455	.705882	.3125	.454545	.294118	1
m	28	32	40	20	20	.583333	.666667	.833333	.621053	.7	.8	.416667	.5	.666667	.416667	.5	.333333	1
a	28	36	39	25	14	.583333	.75	.8125	.663158	.717949	.923077	.520833	.641026	.78125	.291667	.358974	.21875	1
/	25	31	38	18	20	.520833	.645833	.791667	.578947	.657895	.815789	.375	.473684	.642857	.416667	.526316	.357143	1
:	24	33	38	19	19	.5	.6875	.791667	.589474	.631579	.868421	.395833	.5	.666667	.395833	.5	.333333	1
d	37	35	43	29	14	.770833	.729167	.895833	.747368	.860465	.813953	.604167	.674419	.805556	.291667	.325581	.194444	1
s	30	35	40	25	15	.625	.729167	.833333	.673684	.75	.875	.520833	.625	.769231	.3125	.375	.230769	1
o	33	7	33	7	26	.6875	.145833	.6875	.421053	1	.212121	.145833	.212121	.35	.541667	.787879	.65	o
n	39	39	44	34	10	.8125	.8125	.916667	.810526	.886364	.886364	.708333	.772727	.871795	.208333	.227273	.128205	1
a	28	28	39	17	22	.583333	.583333	.8125	.589474	.717949	.717949	.354167	.435897	.607143	.458333	.564103	.392857	o

Vergleicht man die Klassifikation in der Tabelle 7.1. mit der Kategorisierung, die ALTMANN, LEHFELDT (1980, 263) für Transkriptionen des amerikanischen Englisch (aE) auf der Basis einer Studie von ROBERTS aus dem Jahre 1965 vorgenommen haben, so ergibt sich:

- (a) In die gleiche Klasse gehören
 [m e d s n] = aggressiv + attraktiv;
 [j] b g o v] = semiaggressiv + semiattraktiv ("Mittelfeld").
 (b) Um nur eine Klasse verschieden sind (z.B. ist [ʒ] im aE nichtaggressiv + semiattraktiv; im Deutschen nichtaggressiv + nichtattraktiv)
 [z ʒ t a k l e u h i].
 (c) Unterschiedliche statistische Eigenschaften haben
 [t] (aE: semiaggressiv + semiattraktiv; Deutsch: aggressiv + attraktiv);
 [p] (aE: aggressiv + semiattraktiv; Deutsch: nichtaggressiv + nichtattraktiv).

(Es wurden nur solche Symbole verglichen, die für deutsche und amerikanische Aussprache phonetisch adäquat sind).

Für den F- und P-Laut lohnt sich ein Vergleich der einzelnen Parameter, um die Unterschiede spezifizieren zu können:

	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S' _i	S' _i *
F _{aE}	.69	.50	.69	.60	1.	.73	.5	.73	.84	.19	.27	.16
F _D	.70	.66	.82	.68	.85	.80	.54	.66	.79	.28	.34	.21
P _{aE}	.53	.72	.75	.64	.71	.96	.50	.67	.80	.25	.33	.20
P _D	.22	.20	.36	.21	.61	.56	.06	.17	.29	.30	.83	.71

Als unterschiedlich sollen Maßzahlen gelten, deren Differenz > .10 ist. P- und F-Laute unterscheiden sich aus verschiedenen Gründen in beiden Transkriptionen. Die deutsche Symbolverwendung [f] ist assoziativer (As), weil mehr Nachfolgesymbole möglich sind (Ag); bezogen auf die Zahl der Symbole, mit denen F als Vorgänger und/oder Nachfolge Paare bildet (At_i), ist die deutsche Verwendung weniger attraktiv.

Der Gebrauch von [p] ist in beiden Transkriptionen ähnlich

Tabellen 7.1 bis 7.3: Kombinierte assoziative Klassifikation für Transkriptionssymbole bei Aussprache mit schwäbisch-alemannischem und berliner Einfluß (7.1), schwäbisch-alemannischem Einfluß (7.2) und berliner Einfluß (7.3)

7.1	nichtattraktiv	semiattraktiv	attraktiv
nichtaggressiv	3 ˘ + ʀ ˘ ° Y ʒ p y	ŋ u x h	
semiaggressiv	æ o .	k j ʃ ʊ b z g ɛ ɔ ɔ v	i a o ʔ ʔ ç t e . /
aggressiv	° dev	ʔ ʔ	f h : m ə d s a n

7.2	nichtattraktiv	semiattraktiv	attraktiv
nichtaggressiv	X ˘ p y ʒ Y ° æ ŋ z	u x ɔ ʃ ʔ h	
semiaggressiv	o .	k j b ʊ ɛ g a ʔ ʔ ɔ h ʔ v / e ʔ ʔ	: ç t .
aggressiv	° dev	f a	m d ə n s

7.3	nichtattraktiv	semiattraktiv	attraktiv
nichtaggressiv	3 ˘ + ʀ ʀ Y ʒ p k y ŋ o ʔ u h ʒ	j x g . v z	° dev
semiaggressiv	ʔ o .	ɛ a b f ɔ ʔ ʔ t ç h ʔ ʔ ʔ e m / a	
aggressiv		a : s	d n

in der Anzahl der Symbole, mit denen [p] als Vorgänger und/oder als Nachfolger Paare bildet sowie in der Eigenschaft, entweder Vorgänger oder Nachfolger eines anderen Symbols, niemals beides gleichzeitig zu sein (S').

Diese Parameter sind geeignet, die Symbolrelationen innerhalb ein- und desselben Systems zu beschreiben. Phonologisch aufgefaßt ist das System das Phonemsystem; methodisch gesehen das einheitliche Transkriptionssystem. Anhand des Resultats kann noch einmal das wissenssoziologisch gestützte Argument aufgegriffen werden, ob und wie stark der Vollzug von Transkriptionsregeln in das Ergebnis einfließt. Es haben sich bei gleichem Transkriptionssystem, aber unterschiedlichem Material, andere Verwendungseigenschaften herausgestellt. Die Symboldistribution bildet in diesem Beispiel mehr als nur die Kombinatorik des Notationssystems ab; die Regeln des Systems sind weit genug gefaßt, um phonematische oder lautsprachliche Kombinationen abbilden zu können. Sachlich gesehen, d.h. wenn die statistischen Unterschiede der Verwendung von [p] als Ausspracheunterschied interpretiert wird, ist es evident, daß in der deutschen Umgangssprache [p] weit häufiger wie [b] ausgesprochen wird als in der Standardaussprache des aE bzw. als phonologisch überhaupt erfaßt werden kann. Es läßt sich auch das Argument vom Einfluß des Vorwissens auf die Transkription relativieren: Phonetisch erwartbar ist, daß alle stimmlosen Fortis-Klusile umgangssprachlich wie entstimmlichte stimmhafte Lenis-Explosivlaute ausgesprochen werden, aber nur [p] ist auffällig weniger notiert worden als [k] oder [t].

Nicht nur in der distributionellen Eigenschaft eines einzelnen Symbols, sondern auch in der Symbolähnlichkeit läßt sich ein Unterschied zwischen verschiedenen Sprachen (bei Verwendung des gleichen Transkriptionssystems) nachweisen. Zwei Symbole können sich aufgrund ihrer distributionellen Eigenschaften ähnlich sein; wenn sich die Paarähnlichkeit von Material zu Material ändert, kann die Differenz nicht auf Eigenschaften des Transkriptionssystems beruhen. Exemplarisch kann das am Symbolpaar [t] und [v] gezeigt werden, das sich im aE wie im Deutschen klassifikatorisch unterscheidet und zusätzlich durch [t] das aE und das Deutsche um eine Klasse trennt (Berechnung s. Tab. 8). Der Wert von $S_{\text{attraktiv}}$ drückt aus, ob Symbole, die vor [t] stehen, auch vor [v] stehen können (der Parameter $S_{\text{aggressiv}}$ im umgekehrten Sinne; vergleiche

Tabelle 9.1: Distributionelle Frequenz der Transkriptionssymbole in schwäbisch-alemannisch beeinflusste Aussprache

	h	η	æ	~	æ	o	/	ø	:	?	ä	ç	ε	l	ö	r	ſ	o	λ	γ	z	a	b	d	e	f	g	h	i	j	k	l	m	n	o	p	s	t	u	v	x	y	z	+	•	Σ				
h	1	6	1	3	4	11	-	4	2	31	-	5	3	1	3	-	-	2	-	-	-	1	1	4	-	1	-	1	-	4	-	-	12	18	-	5	2	-	9	-	-	-	-	1	11	136				
η	-	-	-	-	-	1	13	-	2	-	-	-	-	-	-	-	-	-	-	-	-	-	-	11	-	6	-	-	-	38	-	1	1	1	1	1	27	-	-	-	-	-	-	-	h	104				
æ	2	1	8	-	1	3	2	1	-	-	1	-	-	1	-	-	-	-	-	-	-	-	4	-	2	-	-	1	-	1	3	-	-	2	4	-	6	-	-	-	-	-	-	-	η	44				
~	10	12	1	8	-	9	3	2	29	12	-	7	7	9	9	10	4	3	5	1	-	13	-	-	7	22	1	14	6	13	-	2	8	7	11	17	-	-	51	8	-	9	1	-	16	-	1	æ	348	
æ	-	-	1	-	-	-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	~	4		
o	-	-	5	-	-	-	-	-	-	-	3	10	3	-	10	2	7	10	1	-	-	5	-	6	-	4	-	-	2	-	-	-	1	3	1	-	1	-	2	-	-	-	-	-	-	-	-	o	78	
/	-	1	5	2	-	-	3	4	-	-	1	-	-	1	-	1	-	-	-	-	-	-	-	-	3	-	1	2	-	-	-	-	-	4	-	-	-	-	-	1	-	-	-	-	-	-	-	/	29	
ø	-	7	-	-	-	-	-	-	-	-	2	1	-	1	-	-	-	-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	11	54	-	-	-	-	-	-	-	-	-	-	-	ø	77		
:	-	1	-	-	-	-	-	-	-	-	1	-	-	1	-	-	-	-	-	-	-	-	-	44	213	-	-	-	2	-	-	-	-	-	1	-	1	-	-	1	56	-	-	5	-	-	-	:	326	
?	2	8	-	11	-	1	1	33	6	-	3	7	24	-	11	4	-	4	2	-	1	-	1	4	23	1	2	4	7	-	-	4	-	7	8	2	-	9	6	-	6	-	1	-	-	-	?	203		
ä	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	1	2	-	-	3	-	-	-	1	-	-	-	-	ä	10		
ç	-	1	3	-	-	1	2	1	-	1	1	-	3	2	2	-	-	-	-	-	-	-	11	-	1	4	4	2	20	3	-	-	-	7	9	7	-	16	1	1	1	1	-	-	-	-	ç	109		
ε	-	1	5	-	-	-	-	-	-	-	1	-	1	3	1	1	-	-	1	-	-	-	29	-	-	30	8	-	2	19	-	-	-	7	14	-	20	-	8	-	-	1	-	-	-	-	-	ε	194	
l	-	-	24	-	-	10	3	1	22	-	5	6	30	-	44	1	9	20	4	-	-	9	-	4	-	1	5	-	-	2	1	-	-	11	6	2	-	-	-	-	1	3	-	-	-	-	l	226		
ö	2	6	-	-	-	1	-	1	14	2	-	1	-	4	-	4	2	-	2	1	-	-	-	8	-	4	5	2	-	2	3	3	1	3	-	3	8	-	5	-	-	-	-	-	-	-	ö	89		
r	1	1	3	2	-	2	2	6	11	3	-	3	4	5	1	-	1	3	36	-	4	-	1	3	7	78	8	-	1	-	3	-	27	5	44	-	46	9	-	6	21	-	-	-	-	-	r	349		
ſ	5	5	-	7	-	2	1	1	29	5	-	5	-	8	-	3	1	3	7	2	-	2	1	-	2	15	-	5	5	8	-	4	3	1	11	4	-	11	10	1	11	-	1	-	-	-	-	ſ	180	
o	-	2	1	-	2	-	-	7	-	-	1	2	-	13	8	5	-	2	-	-	-	1	-	-	4	-	2	3	3	-	2	1	7	3	18	-	3	2	-	-	1	-	-	-	-	-	o	93		
λ	-	3	3	-	-	2	-	10	1	-	5	4	-	2	1	76	-	-	-	-	-	4	-	-	1	6	-	2	5	1	1	2	5	6	9	16	-	17	5	-	1	-	-	-	-	-	λ	188		
γ	1	1	-	5	-	-	1	1	13	1	5	1	2	1	1	-	-	1	1	-	-	-	-	5	9	-	8	-	1	-	-	-	3	10	13	-	10	2	-	4	11	-	1	-	-	-	γ	112		
z	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	z	1		
+	-	1	-	3	-	-	-	1	14	-	1	-	3	2	-	10	-	2	2	2	-	2	-	-	2	7	-	1	-	1	1	-	-	1	1	4	3	2	-	10	1	4	-	-	-	-	-	+	81	
•	1	-	7	2	-	2	-	1	-	-	3	2	-	2	-	-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	-	1	10	28	1	-	11	-	-	-	-	-	-	-	•	72		
Σ	-	-	-	2	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	Σ	3	
h	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	h	11	
η	2	5	3	3	-	2	-	1	13	5	-	1	-	6	2	1	1	1	4	2	-	-	-	-	3	13	1	1	8	2	-	1	2	8	1	6	1	-	3	7	-	2	2	-	-	-	-	η	118	
æ	1	-	-	10	-	3	-	3	3	-	-	6	2	-	3	-	-	4	1	-	1	-	-	-	2	-	3	2	-	7	1	-	4	1	6	1	-	2	1	-	2	1	-	-	-	-	-	æ	70	
~	6	-	2	46	1	5	-	7	13	7	-	4	19	10	1	35	2	13	9	5	-	5	6	-	10	-	4	12	5	2	5	24	1	-	4	7	13	1	-	29	-	8	7	-	-	3	-	-	~	331
o	-	1	1	-	-	3	-	5	9	1	-	2	8	5	8	2	4	2	3	-	-	1	2	-	3	12	1	3	7	-	-	1	2	9	14	29	1	-	23	6	-	-	-	1	-	-	-	o	169	
ç	2	1	-	14	-	4	1	5	12	3	-	2	1	7	-	14	-	1	10	9	-	-	1	-	2	6	-	2	4	-	4	-	1	1	2	7	1	-	3	2	-	1	-	2	2	-	-	-	ç	132
ε	-	2	5	18	-	3	-	6	3	1	1	3	6	2	-	8	-	2	1	3	-	-	1	-	6	-	2	3	1	-	2	1	-	5	2	7	2	-	4	2	-	-	-	-	-	-	-	ε	102	
l	9	3	-	7	-	2	2	9	16	-	6	8	9	-	26	3	-	5	1	-	-	3	-	3	2	7	-	2	1	6	2	2	-	-	3	8	2	-	2	3	-	-	-	-	-	-	-	l	154	
ö	1	3	-	5	-	2	-	9	2	-	-	1	3	7	6	2	1	1	-	-	1	1	-	1	4	2	-	2	4	-	-	1	1	1	4	2	1	1	2	4	-	2	-	-	3	-	-	ö	80	
r	2	1	-	1	-	-	1	2	2	-	8	8	-	1	9	1	6	-	-	-	-	5	-	18	-	2	4	-	-	-	-	-	-	-	2	-	4	-	-	-	-	-	-	-	-	-	-	r	80	
ſ	-	-	1	-	3	4	2	1	-	-	2	2	1	-	7	1	1	1	6	-	-	1	-	-	-	3	1	-	1	-	-	-	-	1	2	3	2	-	6	-	-	1	-	1	-	-	-	ſ	54	
o	-	2	-	6	-	-	-	14	1	-	-	7	2	-	11	7	3	9	3	-	-	1	-	2	1	15	5	2	1	1	2	-	-	2	5	5	-	6	8	-	2	-	9	-	-	-	-	o	132	
λ	10	2	-	16	-	-	1	4	18	6	-	7	8	7	-	30	1	2	12	2	-	5	12	-	7	6	11	3	1	1	5	3	3	-	4	7	6	2	1	9	5	-	1	1	-	1	-	-	λ	220
γ	7	16	-	26	-	1	2	7	42	10	1	8	9	25	-	18	5	4	29	5	-	12	10	1	-	7	3	46	4	7	5	19	1	8	8	4	17	9	1	-	12	28	1	13	-	5	-	1	γ	437
z	7	2	-	4	-	1	-	2	14	5	-	4	-	8	3	2	2	2	4	-	-	1	3	-	-	2	8	-	5	4	2	-	1	-	3	6	13	1	-	1	2	-	12	1	-	3	-	-	z	129
+	-	-	-	1	-	-	-	-	-	-	-	1	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	+	8	
•	11	10	1	15	-	6	2	6	57	27	1	4	15	27	-	25	2	5	19	8	-	4	6	-	4	6	30	4	11	11	5	8	14	5	3	18	26	30	1	7	23	23	20	-	-	-	-	•	500	
h	2	5	-	16	-	5	-	1	5	8	-	6	9	7	-	10	3	4	7	-	-	4	5	-	5	-	1	6	1	2	7	3	3	2	1	1	1	3	-	48	3	6	4	1	2	-	-	-	h	198
η	-	1	-	1	-	-	-																																											

Tabelle 9.2: Distributionelle Frequenz der Transkriptionssymbole für berlinisch beeinflusste Aussprache

	ll	h	η	ə	ʊ	æ	.	•	/	ø	:	?	ɒ	ɑ	ɘ	ɛ	ɪ	ɔ	ʀ	ʃ	ɔ	χ	ʏ
ll	-	1	-	-	-	-	-	6	-	-	3	1	3	1	4	1	-	4	-	-	1	1	-
h	1	-	-	-	-	-	-	3	-	-	1	1	-	-	-	-	1	-	-	-	-	-	-
η	1	-	1	1	-	-	-	2	2	-	-	-	-	-	-	-	-	-	-	-	1	-	-
ə	3	2	1	1	-	6	1	2	8	4	-	3	4	2	15	6	1	1	1	-	10	-	-
ʊ	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
æ	-	-	-	3	-	-	-	-	1	-	-	4	4	-	-	1	-	4	4	3	-	-	-
.	-	-	-	-	-	-	-	1	-	-	-	-	-	-	2	-	-	-	-	-	-	-	-
•	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
/	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
ø	1	3	-	3	-	-	-	25	5	-	6	1	14	-	10	2	4	-	-	1	-	-	1
:	1	2	-	2	-	-	-	4	2	-	-	-	-	-	-	1	-	-	-	-	-	-	-
?	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
ɒ	-	-	-	6	-	-	-	2	-	3	2	1	-	17	1	10	10	1	-	-	-	-	-
ɑ	5	3	-	-	-	-	2	9	13	-	2	2	3	-	6	-	3	-	2	-	1	1	-
ɘ	-	1	-	-	-	-	3	9	2	1	7	1	2	2	-	1	1	1	14	-	-	-	-
ɛ	2	1	-	1	-	2	-	4	9	3	-	2	-	4	-	1	1	1	5	-	1	2	-
ɪ	2	-	-	-	-	1	2	3	-	2	1	4	1	3	6	2	2	-	-	1	2	-	-
ɔ	-	-	4	-	-	-	3	7	1	-	4	-	-	5	-	30	-	-	-	-	-	-	-
ʀ	-	-	1	-	-	-	1	7	1	3	3	-	-	1	-	-	1	-	-	-	-	-	-
ʃ	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
ɔ	-	-	1	-	-	-	-	2	3	-	1	1	1	-	1	-	2	-	1	1	-	-	-
χ	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
ʏ	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
ɜ	-	-	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
ʌ	8	-	-	-	-	-	1	9	3	1	-	1	1	3	1	-	-	2	-	-	1	-	-
ɑ	1	-	-	11	-	-	1	3	1	-	1	2	-	1	6	1	2	3	-	-	1	-	-
b	7	-	-	19	-	1	-	2	7	4	-	20	10	1	2	27	5	1	4	4	1	1	1
d	-	-	1	1	-	-	-	1	3	2	-	3	1	4	24	3	4	-	-	1	-	-	-
e	1	1	-	2	-	-	-	1	2	3	-	3	5	1	-	5	-	3	4	4	-	-	1
f	-	-	-	15	-	2	-	2	1	-	-	1	2	-	-	4	-	-	1	1	-	-	-
g	6	2	-	7	2	1	-	-	14	-	10	4	9	-	11	1	4	-	-	-	-	-	-
h	1	3	-	2	-	1	-	3	4	-	6	2	4	4	7	-	2	2	1	-	-	2	-
i	-	-	2	-	-	-	1	3	-	-	5	3	-	-	1	-	5	-	-	-	-	1	-
j	-	-	1	1	-	-	1	1	-	-	-	-	-	-	2	-	2	-	-	-	-	-	-
k	2	-	3	-	-	-	8	3	-	4	4	2	-	8	2	-	4	2	-	1	-	-	-
l	19	1	-	4	-	-	8	5	-	7	2	-	-	16	2	1	9	2	-	-	5	-	1
m	7	1	-	11	-	3	-	3	19	10	1	13	4	10	-	14	1	4	9	2	1	2	-
n	5	-	2	1	-	-	3	4	4	-	5	2	4	2	5	-	-	-	-	-	1	-	-
o	-	-	-	-	5	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
p	5	5	-	5	-	6	25	18	1	12	7	8	-	8	4	3	11	1	-	-	1	-	1
s	1	1	-	9	-	4	2	1	-	3	3	-	-	3	4	1	3	-	-	-	4	-	-
t	1	-	-	-	-	-	1	-	-	-	1	-	5	-	-	-	2	-	-	-	-	-	-
u	-	-	1	-	-	-	-	-	-	5	5	-	2	23	-	-	5	4	-	-	-	-	-
v	-	-	2	-	-	-	4	7	2	-	-	1	1	-	2	-	2	-	1	-	-	-	-
x	-	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
y	-	-	4	-	-	-	3	2	1	-	3	8	-	7	-	1	5	4	-	-	-	-	-
z	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
+	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
•	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Σ	85	16	2	3	198	7	84	73	70	105	5	17	5	78	276	61	98	32	91	249	15	77	2

	a	b	d	e	f	g	h	i	j	k	l	m	n	o	p	s	t	u	v	x	y	z	+	•	Σ	
-	1	6	1	1	1	8	-	7	-	2	13	13	-	-	6	1	-	5	-	1	-	-	-	ll	92	
-	-	4	1	-	1	-	-	1	3	-	-	2	-	-	-	6	1	-	-	1	-	-	-	h	27	
-	-	1	-	-	3	-	-	1	1	-	2	1	-	-	2	-	-	-	-	-	-	-	-	η	19	
1	1	9	2	5	2	6	-	1	2	2	3	4	-	-	10	4	-	5	1	1	6	-	-	e	136	
2	-	-	-	-	-	-	-	-	-	-	-	-	-	2	-	-	-	1	-	-	-	-	-	•	2	
6	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	æ	34	
-	-	-	-	-	-	-	-	-	-	-	-	2	-	-	-	-	-	-	-	-	-	-	-	.	5	
-	-	-	-	-	-	1	-	-	-	-	-	12	37	-	-	-	-	-	-	-	-	-	-	-	•	52
-	22	84	1	-	12	-	-	-	-	-	-	1	-	-	1	-	-	25	-	-	19	-	-	•	165	
-	6	20	1	3	5	6	-	1	2	2	7	9	1	5	4	7	-	3	1	-	4	-	-	/	163	
-	-	2	-	-	-	-	-	-	-	-	-	1	3	-	-	-	-	-	-	-	-	-	-	ø	11	
9	-	2	11	3	-	-	11	27	-	1	1	19	9	12	-	11	1	4	-	-	4	-	-	:	144	
17	-	1	10	4	-	-	-	11	-	-	-	3	14	-	1	-	3	-	-	6	-	-	-	•	77	
1	-	1	4	-	-	2	2	-	-	-	-	3	3	1	-	-	-	-	1	1	-	-	-	?	72	
-	2	9	1	4	1	2	-	1	2	3	7	8	1	-	5	1	-	6	-	1	3	-	-	•	109	
-	4	7	36	1	1	-	-	1	-	26	4	13	-	-	41	2	-	3	13	-	2	-	-	•	200	
-	-	6	1	5	-	2	-	-	2	-	4	7	1	-	6	7	1	6	-	-	-	-	-	•	89	
-	-	2	-	2	-	2	-	-	-	6	3	7	1	-	4	2	-	-	-	-	1	-	-	•	61	
-	1	5	-	1	1	1	-	1	-	5	4	12	-	-	10	5	-	1	-	-	-	-	-	•	101	
-	7	4	-	5	-	-	-	-	-	1	7	3	1	-	-	4	2	-	-	8	-	1	-	•	61	
-	-	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	•	5	
-	1	1	1	-	-	-	-	-	-	-	-	-	1	2	5	-	3	-	1	-	-	-	-	•	29	
-	-	1	-	-	-	-	-	-	1	-	-	2	6	-	-	3	-	-	-	-	-	-	-	•	19	
-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	2	-	-	-	-	-	-	-	•	5	
-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	•	1	
-	4	5	-	1	2	3	-	2	-	2	4	3	-	-	4	4	-	3	1	-	2	-	-	•	71	
-	-	2	-	-	-	3	5	-	-	3	6	1	-	-	4	-	-	-	-	1	-	-	-	•	59	
1	5	1	3	16	-	1	2	24	3	1	4	4	4	1	-	6	-	8	5	-	-	4	-	•	210	
-	1	2	4	1	1	2	1	1	1	-	5	9	21	-	-	4	-	-	1	-	1	2	-	•	105	
-	-	1	-	1	2	3	6	-	-	2	-	2	2	-	-	4	2	-	1	1	4	-	-	•	67	
3	-	1	4	-	-	-	-	-	-	2	1	2	-	-	1	-	-	-	-	-	-	-	-	•	43	
6	-	3	1	1	-	9	1	1	-	7	12	-	-	-	2	1	-	-	-	-	1	1	-	•	117	
-	-	2	-	4	1	-	1	1	1	6	4	-	2	-	4	3	-	1	-	3	5	-	-	•	82	
8	1	-	-	-	-	-	1	-	1	-	2	1	-	-	-	-	-	-	-	-	-	-	-	•	35	
-	-	-	-	-	-	-	1	-	-	1	-	1	-	-	-	-	-	-	-	-	-	-	-	•	13	
-	-	5	9	5	-	1	1	-	1	-	2	2	4	5	-	14	6	-	2	1	-	9	-	•	110	
-	6	4	3	2	2	3	5	1	2	1	1	18	5	-	1	4	4	-	3	-	-	1	-	•	148	
-	4	7	27	6	3	3	20	-	1	1	1	8	7	2	1	6	5	3	6	-	1	8	1	•	237	
-	1	4	5	1	6	3	2	-	-	4	2	11	-	-	2	1	-	1	1	1	1	-	-	•	89	
-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1	-	-	-	-	3	-	-	-	•	10	
-	2	30	7	4	7	3	5	2	2	1	5	14	18	2	2	8	3	7	-	-	3	-	-	•	246	
-	1	1	4	3	1	-	1	2	2	-	1	-	5	-	-	15	-	3	-	-	-	-	-	•	79	
-	-	1	3	1	2	2	-	-	-	-	-	2	-	-	-	-	-	-	-	-	-	1	-	•	23	
3	-	1	3	-	-	-	-	13	-	-	-	1	-	-	-	1	-	-	-	-	-	-	-	•	67	
-	1	-	4	-	1	1	1	-	-	-	3	6	-	-	-	1	1	1	2	-	-	1	-	•	46	
-	1	2	1	-	-	1	-	-	-	1	-	-	-	-	-	-	-	-	3	-	-	-	-	•	12	
-	2	-	2	1	-	-	-	-	-	-	-	1	-	-	3	24	1	-	-	-	-	-	-	•	72	
-	-	-	-	-	-	-	1	-	-	-	-	-	1	-	-	-	-	-	-	-	-	-	-	•	2	
-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	•	-	
1	79	121	56	100	20	160	91	184	28	31	77	-	3622	78	276	61	98	32	91	249	15	77	90	25	2	3622

[illegible]

! " # \$ % & ' () * + , - . / : ; { | } ~ ¢ £ ¤ ¥ ¦ § ¨ © ª « ¬ ® ¯ ° ± ² ³ ´ µ ¶ · ¸ ¹ º » ¼ ½ ¾ ¿ À Á Â Ã Ä Å Æ Ç È É Ê Ë Ì Í Î Ï Ñ Ò Ó Ô Õ Ö × Ø Ù Ú Û Ü Ý Þ ß à á â ã

10.2 (Fortsetzung)

[illegible]

11 h q e ~ æ . / ø : ? p æ ç ε ι ρ ρ ∫ q χ γ ζ α b d e f g h i j k l m n o p s . t u v x y z +

Tabelle 11: Zusammenstellung der statistischen Klassifikationen P (bevorzugt) und I (unzulässig) gemäß Tab. 10.1

vor * steht bevorzugt:	* nach * folgt bevorzugt:	vor * ist unzulässig:	* nach * ist unzulässig:
1. Vokale			
[: · b d f v] → [i]	[h o f i y z]	[ɔ s a a o o] → [i]	[e]
[ʊ ʔ f i m] → [u]	[n ɕ i n s]	[· i j] → [u]	[ʰ ə · e i]
[: · a d] → [e]	[ɔ g i m n s]	[ɔ b f i] → [e]	[ʰ]
[ʊ ʔ d j v] → [ɛ]	[ʊ ɕ x i h]	[· u] → [ɛ]	[a e i v]
[g t v z] → [ə]	[ʰ h ʊ k ʰ]	[· o u] → [ə]	[o u o]
[ʊ j m v z] → [a]	[t g i b]	[ʊ e t ɔ] → [a]	[a i]
[ʊ ʔ f d f h] → [a]	[ɔ e i]	[h a o o] → [a]	[a i k o u]
[: · i s] → [o]	[g n v]	[ʰ ə a u ɔ] → [o]	[o i]
[a ʒ f k] → [ɔ]	[ɔ ɸ b f]	[ʰ i · ɔ] → [ɔ]	[a ɕ a e i]
[: · d s t] → [u]	[ɔ f h x z]	[ə / a m] → [u]	[s]
[ʊ ʔ j m t] → [a]	[n m n s]	[ə · /] → [o]	[ʔ a]
[ɔ] → [ɸ]	[· n t z]	[ɸ] →	

Tabelle 11 (1. Fortsetzung)

vor * steht bevorzugt:	* nach * folgt bevorzugt:	vor * ist unzulässig:	* nach * ist unzulässig:
[ʰ k ʔ] → [ɛ]	[n]	[ɛ] →	
[: · ʔ f i] → [y]	[ɔ v]	[y] →	
[f] → [Y]	[s z]	[Y] →	
[ə e e i u y] → [ɔ]	[ʰ h ɔ /]	[ɔ ʔ ɕ h i] → [ɔ]	[n a i]
2. Klusile			
[ɔ a] → [ɔ]	[ə · i i]	[: · ʔ v] → [b]	[/ e]
[h · / i n] → [d]	[ə : a]	[k] → [d]	[h t x]
[a e o s] → [g]	[n ə ʊ i]	[· ʔ v] → [g]	[ɕ h v]
[/ j] → [p]	[ʊ v]	[p] →	
[a i n s x] → [t]	[ə ʊ · o s]	[· ʔ d j] → [t]	[ɔ ɔ ʒ]
[h ə ɔ] → [k]	[ʊ ə a ɔ]	[· ʔ a] → [k]	[d]
[h n o s] → [ʔ]	[ə ʊ / ʔ a]	[ʰ o ʒ j] → [ʔ]	[ʰ ʊ b ʔ /]
3. Nasale			
[ɔ o e m] → [m]	[ʰ : a i]	[· v] → [m]	[ʊ ɔ u]

Tabelle 11 (2. Fortsetzung)

vor * steht bevorzugt:	*	nach * folgt bevorzugt:	vor * ist unzulässig:	*	nach * ist unzulässig:
$\begin{bmatrix} \text{t} & \text{d} & \text{e} & \text{o} & \text{x} & \text{z} \\ \text{e} & \text{a} & \text{ø} & & & \text{x} \end{bmatrix} \Rightarrow [\text{n}] \Leftarrow \begin{bmatrix} \text{d} & \text{h} & \text{t} \\ \text{o} & \text{h} & \text{t} \end{bmatrix}$					$\Leftarrow [\text{n}] \Rightarrow [\text{ŋ} \text{ t x}]$
$[\text{x} \text{ , g t o}] \Rightarrow [\text{ŋ}] \Leftarrow [\text{ə v g}]$					$[\text{n}] \Leftarrow [\text{ŋ}] \Rightarrow$
4. Lateral					
$\begin{bmatrix} \text{a} & \text{b} & \text{e} & \text{g} & \text{i} \\ \text{o} & \text{t} & \text{e} & \text{a} \end{bmatrix} \Rightarrow [\text{l}] \Leftarrow \begin{bmatrix} \text{a} & \text{g} & \text{t} \\ \text{o} & \text{d} & \text{s} \end{bmatrix}$			$[\text{.} \text{ , / v}] \Leftarrow [\text{l}] \Rightarrow [\text{,} \text{ , r}]$		
5. Vibranten					
$\Rightarrow [\text{R}] \Leftarrow$					$\Leftarrow [\text{R}] \Rightarrow$
$[\text{ə} \text{ e g k p t}] \Rightarrow [\text{v}] \Leftarrow [\text{:} \text{ , e t o a}]$					$[\text{.} \text{ / m}] \Leftarrow [\text{v}] \Rightarrow [\text{/ t v}]$
6. Konstruktive					
$\begin{bmatrix} \text{ç} & \text{o} & \text{s} & \text{y} \\ \text{o} & & & \text{y} \end{bmatrix} \Rightarrow [\text{v}] \Leftarrow \begin{bmatrix} \text{o} & \text{a} & \text{t} \\ \text{e} & \text{a} & \text{e} \end{bmatrix}$			$\begin{bmatrix} \text{e} & \text{g} & \text{h} & \text{j} & \text{z} \\ \text{h} & & & & \text{r} \end{bmatrix} \Leftarrow [\text{v}] \Rightarrow \begin{bmatrix} \text{t} & \text{d} & \text{g} \\ \text{h} & \text{l} & \text{m} \end{bmatrix}$		
$\begin{bmatrix} \text{y} & \text{i} & \text{l} & \text{u} \\ \text{ə} & & & \text{v} \end{bmatrix} \Rightarrow [\text{z}] \Leftarrow [\text{ə} \text{ , t a u}]$		$[\text{t} \text{ , ? t t}] \Leftarrow [\text{z}] \Rightarrow [\text{r ç h v}]$			
$\Rightarrow [\text{z}] \Leftarrow [\text{ɔ}]$					$\Leftarrow [\text{z}] \Rightarrow$
$[\text{t} \text{ x s}] \Rightarrow [\text{j}] \Leftarrow [\text{:} \text{ , e o a e}]$					$[\text{.} \text{ ,}] \Leftarrow [\text{j}] \Rightarrow [\text{? t s t v}]$
$\begin{bmatrix} \text{ɔ} & \text{i} & \text{o} & \text{u} \\ \text{ɔ} & \text{ə} & \text{ç} \end{bmatrix} \Rightarrow [\text{f}] \Leftarrow \begin{bmatrix} \text{ə} & \text{a} \\ \text{ɔ} & \text{y i y} \end{bmatrix}$					$[\text{.} \text{ ,}] \Leftarrow [\text{f}] \Rightarrow [\text{ç a e}]$

abelle 11 (3. Fortsetzung)

or * steht bevorzugt:	* nach * folgt bevorzugt:	vor * ist unzulässig:	* nach * ist unzulässig:
$[o \circ y e i t] \rightarrow [s] \leftarrow [h \text{ } j \text{ } / \text{ } i \text{ } u]$		$[, \text{ } / \text{ } u \text{ } j] \leftarrow [s] \Rightarrow [o \text{ } x]$	
$[a \text{ } n] \rightarrow [j] \leftarrow [a \text{ } p \text{ } t]$		$[o] \leftarrow [j] \Rightarrow [s]$	
$[e \text{ } i \text{ } l] \rightarrow [q] \leftarrow [o \text{ } h \text{ } m \text{ } t \text{ } v]$		$[o \text{ } o \text{ } f \text{ } g \text{ } z] \leftarrow [q] \Rightarrow [o \text{ } o \text{ } a \text{ } i]$	
$[a \text{ } o \text{ } p \text{ } u] \rightarrow [x] \leftarrow [, \text{ } o \text{ } / \text{ } n \text{ } t]$		$[o \text{ } d \text{ } n \text{ } s] \leftarrow [x] \Rightarrow$	
$[e] \rightarrow [x] \leftarrow [h \text{ } j]$		$\leftarrow [x] \Rightarrow$	
$[i \text{ } q \text{ } h \text{ } n \text{ } u] \rightarrow [h] \leftarrow [i \text{ } / \text{ } a \text{ } h]$		$[g \text{ } i \text{ } v \text{ } z] \leftarrow [h] \Rightarrow [o \text{ } i \text{ } v]$	
7. Pausen			
$[a \text{ } o \text{ } a \text{ } h \text{ } m \text{ } o] \rightarrow [i] \leftarrow [x \text{ } ? \text{ } h \text{ } j]$		$[i \text{ } e \text{ } ?] \leftarrow [i] \Rightarrow [i \text{ } / \text{ } o \text{ } i \text{ } o \text{ } a]$	
$[i \text{ } o \text{ } h \text{ } s \text{ } x] \rightarrow [/] \leftarrow [h \text{ } o \text{ } / \text{ } d \text{ } h \text{ } p]$		$[h \text{ } o \text{ } b] \leftarrow [/] \Rightarrow [h \text{ } o \text{ } i \text{ } u]$	
8. Sonderzeichen			
$[i \text{ } n \text{ } x \text{ } /] \rightarrow [h] \leftarrow [g \text{ } d \text{ } o \text{ } k \text{ } p]$		$[? \text{ } d] \leftarrow [h] \Rightarrow [? \text{ } / \text{ } v \text{ } a \text{ } a]$	
$[h] \rightarrow [\sim] \leftarrow [a]$		$\leftarrow [\sim] \Rightarrow$	

Tabelle 11 (4. Fortsetzung)

vor * steht bevorzugt:	* nach * folgt bevorzugt:	vor * ist unzulässig:	* nach * ist unzulässig:
[ɐ d h j m z] ⇒ [:] ⇐ [æ h i m]			⇐ [:] ⇒ [ɔ b z]
[^h j l t z] ⇒ [·] ⇐ [æ e f i]			[· · ɔ] ⇐ [·] ⇒ [¹¹ æ, o, tv / ɔ b g l m z]
[æ ɸ b f g x] ⇒ [·] ⇐ [n m ŋ]			[· · ɪ] ⇐ [·] ⇒ [æ, o / t e s t]
	⇒ [·] ⇐ [h]		⇐ [·] ⇒
[j v] ⇒ [°] ⇐ [a a]			⇐ [°] ⇒
[^h / ɐ ɔ ɔ f] ⇒ [·] ⇐ [b d g v z]			[· · o] ⇐ [·] ⇒ [¹¹ æ, o, tv / ɔ b g l m z]

In der Tabelle 12 werden die schwäbisch-alemannisch und berlinisch beeinflusste Aussprache untereinander und mit dem beiden gemeinsamen Klassifikationsfeld verglichen; jedes Symbol, das zur bevorzugten oder unzulässigen Kombination gehört, wird einzeln aufgeführt. Um einen globaleren Überblick zu haben, wird in Tabelle 13 die statistische Relation von Lautgruppen angegeben. Die Lautgruppen sind so zusammengestellt und summiert worden, daß ein Vergleich mit der Klassifikation von ALTMANN, LEHFELDT (1980, 316ff.) des Materials von LINDNER (1975) möglich ist.

Das Material von LINDNER bestand in Schrifttexten, so daß mit ihm der Unterschied zur Umgangssprache herausgestellt werden kann. Die vorliegenden Transkriptionen wurden auf 36 Symbole reduziert; Sonderzeichen, Pausen, [ə], [ɔ], [ʔ] und [χ] wurden nicht berücksichtigt. Diphtonge sind im Gegensatz zur LINDNERSchen Auflistung nicht gesondert ausgewertet worden, sie sind in den Vokalkombinationen enthalten.

In 20 von 64 möglichen Kombinationen weicht die Transkription der Umgangsaussprache von der schriftsprachlichen Kombinatorik ab.

- Kurze Vokale unterscheiden sich nicht vom normalsprachlichen distributionellen Gebrauch; lange Vokale gehen bevorzugt (⇔ aktuell) Nasalen voran und aktuell (⇔ bevorzugt) stimmlosen Verschlusslauten.

- Stimmhafte Verschlusslaute können mit sich selbst und mit stimmlosen Frikativen verbunden werden, wenngleich nur marginal (⇔ unzulässig), worin sich das in der Umgangssprache gewöhnliche Auslassen von Zwischenvokalen zeigt. Mit stimmhaften Frikativen sind sie aktuell (⇔ marginal), mit Liquiden bevorzugt (⇔ marginal) gekoppelt; in beiden Fällen ein Anstieg der Kombinationsmöglichkeit.

- Stimmlose Verschlusslaute stehen nur aktuell (⇔ bevorzugt) vor langen oder kurzen Vokalen; sie verbinden sich bevorzugt mit Liquiden (⇔ marginal), aber nur marginal mit stimmhaften Reibelauten (⇔ aktuell).

- Die Kombination der stimmhaften Reibelauten weicht stark vom Schrifttext ab; sie treten vor stimmhaften Verschlusslauten, stimmlosen Verschlusslauten, stimmlosen Reibelauten, Nasalen und Liquiden marginal auf (⇔ unzulässig in allen Fällen).

- Stimmlose Frikative traten weniger häufig vor langen Vokalen (aktuell ⇔ bevorzugt), aber bevorzugt vor stimmhaften Frikativen

Tabelle 12 (2. Fortsetzung)

* bevorzugt vor *	nach *	unzulässig vor *	nach *
2. Klusile			
b ug o:n sa o:n-af be o:n-----y	a, i --i a--i-æm	:?v --?-dt :-----	øæ -- --
d ug h:/in sa --/in° be -/-n-s	a:æei æ:--e-fus æ:a-eiu-	k kæ·o :? :?-æ :-----	htx ht-bk -t---f çhv --v --
g ug æosm ^h sa æ--n-h-i be -----η-o-u	ηæi ^h ηæ-i ^h -æi-----e		
p ug /f sa -f be /f	hv -- x-x		
t ug a msx ^h øçf sa - sx ^h øçf ^h nn be a ---h-çf---i/	æ'os æ--sy æ--qs-,uç	:?dj ^h :?d ^h :-----kv :-----m	oç -o-b -----m
k ug hen sa h--X be h---	ææo ææo -æ--æ	:?a :-----dm :-----	d dt --
? ug hnos// sa -n-s// ^h be hn---/ --ε	æh/ææioy æ-/ææ-æyæi æ---æ-----i	ioçj ^h l--j ^h :-----m; :-----m;	l h b f v g k i s t z l h b f v g k i s t z -----s--°

Tabelle 12 (3. Fortsetzung)

* bevorzugt vor *	nach *	unzulässig vor *	nach *
3. Nasale			
m ug oæmç:-- sa oæmç:-- be --em:--	l:æioam l:æio--fæ l:æio----	:?v :? :-----tz	øæu --k -----zo
n ug loæx7ææ ^h sa -æo--æa- be l-æx--ææ ^h	f d h t ^h f d h t ^h /æo -dh- ^h /---	j -f n -- --	oæx -æx -o-i
η ug ægio sa æg-æa be -l--x	ævg æv-t. --g--		
4. Lateral			
l ug abegioæa sa ab-g---æau be -begioæa-o	°æioæstz °-çl-d-zu -----stz-	:?v :? :-----	øv -- --i
5. Vibranten			
r ug - sa - be f	b - -		
ʀ ug ægkpt sa æ-k-tuf be æ-g-p---y	:æioæ -æi-æo :æi-æo	°/m :? :-----	øtv --d --

(↔ aktuell) und aktuell vor Nasalen (↔ marginal) auf.

- Nasale als Vorgänger aller anderen Lautgruppen weichen in der Häufigkeit nicht vom Schriftstandard ab.
- Liquide werden weniger oft mit stimmhaften und stimmlosen Verschlusslauten (aktuell ↔ bevorzugt) und stimmlosen Reibelauten (marginal ↔ aktuell) nachfolgend verbunden.

Insgesamt nehmen die Kombinationen stimmhafter Konsonanten mit anderen Lautgruppen zu, was man in der älteren Literatur als "Vorherrschen des stimmlichen Elementes" bezeichnet hätte. Der phonatorisch-artikulatorische Impetus ist mehr legato als staccato, eher fließend als abrupt; der Artikulationsstil ist verbindlicher, "intimer" und vermeidet "hart" klingende Verbindungen. So werden stimmlose Verschlusslaute (bis auf die Kopplung mit den "fließenden" Lauten, d.h. Liquiden) weniger oft kombiniert; die Liquide selbst werden ebenfalls reduziert (Verschleifungen der R- und L-Laute) und der normalsprachliche Gebrauch von stimmlosen Frikativen zugunsten stimmhafter Frikative und Nasale herabgesetzt.

2.3 INDIVIDUELLE ARTIKULATIONSEIGENHEITEN ZWEIER SPRECHER WÄHREND EINES DYADISCHEN GESPRÄCHS IM COLLOQUIAL STYLE

Mit der Auswertung phonetischer Transkriptionen ist - im Kontrast zu rein phonologischen Statistiken - die Möglichkeit gegeben, auch individuelle, paralinguistische Artikulationsbesonderheiten statistisch zu beschreiben. Tabelle 12 enthält bereits einige solcher Unterschiede, seien sie durch dialektale, umgangssprachliche oder individuelle Komponenten bedingt; jedoch sind nur die statistisch bevorzugten oder unzulässigen Distributionen aufgeführt. ALTMANN, LEHFELDT (1980, 266ff.) schlagen einige Verfahren zum statistischen Vergleich verschiedener Sprachen vor, die - in Auswahl - auf einen Vergleich zweier Sprecher angewendet werden sollen.

Sprecherspezifische Artikulationsstile machen einen Teil der paralinguistischen Informationen aus, die dem Hörer den typischen "sound" eines Sprechers vermitteln. Der Höreindruck, soweit er sich an der Aussprache festmachen läßt, kann mit dem IPA-System notiert werden; der Satz von POYATOS (1975), daß alles paralinguistisch

genannt wird, was nicht mit IPA-Symbolen transkribierbar ist (295/296), ist nur ein ironisches Aperçu (zur Notation paraphonetischer features s. WINKLER 1979 und 1982).

Die statistische Analyse soll auf solche Symbole beschränkt werden, die der Transkribent möglicherweise paralinguistisch gemeint hat. Die Selektion ist ein interpretativer Eingriff, der m.E. berechtigt ist, weil ihm eine vollständige statistische Beschreibung vorangegangen ist. In den Tabellen 5 und 7.2/7.3 sind eine Reihe von Symbolen durch die Assoziativitäts-Parameter unterschieden; dort sind sie als Indiz für dialektale und paraphonetische Beschreibungen gewertet worden. Aus diesen klassifikatorisch verschiedenen Symbolen (das sind: [χ ʒ + R K o [n h b i o z · j g . v o f v t ç m a e : s]), den nicht unterschiedlich verwendeten Zeichen [/ ' ~ °] und dem Glottiseinsatz werden anhand des folgenden Fragekataloges 16 Symbole ausgewählt:

- Sind unterschiedliche Häufigkeiten und Kombinationen von Symbolen entstanden, mit denen der Transkribent die Sprecher in Bezug auf
- den Gebrauch des Glottiseinsatzes [?],
 - die Behauchung [h],
 - die Entstimmlichungstendenz [°],
 - Rundungen und Nasalierungen von Vokalen [° ~],
 - die silbische Realisierung [,],
 - die Länge von Lauten [: ·],
 - inverse Atmung [+],
 - den Gebrauch des schwachtonigen E-Lautes [e],
 - die Allophone des R-Lautes bzw. dessen Ersatzrealisation [R b v χ],
 - Stauungen und Pausen [/ ']
- charakterisiert hat?

Möglicherweise überlagern sich dialektaler und paralinguistischer Gebrauch von [~ ° χ v h], der statistisch nicht zu trennen ist. Zumindest gehen die Verwendungseigenschaften von [e . . : ·], vielleicht auch von [h], eindeutiger auf den colloquial style des Gesprächs zurück.

Zur Unterscheidung von den vorangehenden Tabellierungen werden die Sprecher mit "SIE" (vorher: schwäbisch-alemannische Aussprache) und "ER" (vorher: berlinische Aussprache) gekennzeichnet.

Die statistischen Parameter werden nunmehr mit reduziertem K (=16) berechnet, die Kombinatorik jedoch mit K = 50, wobei die Symbole zu Lautgruppen zusammengefaßt sind (ähnlich Tab. 13).

Tabelle 14: Häufigkeitsparameter für 16 Transkriptionssymbole, Divergenz und Korrelation zweier Sprecher

Symbol	SIE f_i	$\hat{p}_i$	ER f_i	$\hat{p}_i$
R	-	-	4	.00439
+	-	-	2	.00219
λ	2	.00128	-	-
~	3	.00193	2	.00219
o	12	.00770	-	-
o	60	.03851	72	.07912
o	63	.04044	25	.02747
.	75	.04814	52	.05714
	82	.05263	76	.08352
h	88	.05648	21	.02308
:	100	.06418	134	.14725
/	126	.08087	109	.11978
.	154	.09884	77	.08462
z	159	.10205	63	.06923
e	258	.16559	109	.11978
o dev	376	.24134	164	.18022

N	1558	910
Σf^2	310632	93266
χ^2	1632.059	729.842
R	.12797	.11263
D	.0698	.0535
V	1.0235	.8956
H ₀	4	4
H ₁	8.2444	8.4215
h	.8111	.8554
Re	.1889	.1446
I _L	.1055	.0933

$r = .838$
 $D = .1399$
 $\chi^2 = 139.85$ (G = 145.58)

Tabelle 14 enthält die Häufigkeitsangaben für beide Sprecher; die Verwendung der 16 Symbole korreliert zwischen beiden Stichproben recht hoch ($r = .838$; Signifikanz $t = 5.75$ bei Tafelwert $t_{.05;14} = 2.14$). So ist die Divergenz mit .14 nicht sehr groß (D kann sich zwischen 0 und 1.41 bewegen; Berechnung aus $\sqrt{\Sigma(\hat{p}_{x_i} - \hat{p}_{y_i})^2}$). Die Stichproben sind jedoch signifikant inhomogen.

In den distributionellen Eigenschaften unterscheiden sich beide Symbolverwendungen (gemäß Tabellen 6.2 und 6.3):

Tabelle 15: Differenzen der Distributionsmaße beider Sprecher

Sy	At	Ag	As	As*	At _i	Ag _i	S	S _i	S _i *	S'	S _i '	S _i '*	$\sqrt{\Sigma(a-b)^2}$
h	.31	.01	.31	.15	.16	.24	.02	.09	.14	.32	.09	.14	.68
dv	.46	.64	.14	.09	.72	.74	.04	.02	.03	.10	.02	.02	1.31
o	.09	.21	.28	.15	.31	.77	.02	.08	.14	.26	.92	.86	1.59
~	.07	.04	.11	.01	.07	.07	.00	.00	.00	.11	.00	.00	.20
.	.20	.60	.18	.20	.50	.79	.16	.06	.42	.03	.28	.42	1.34
:	.27	.09	.04	.10	.29	.17	.13	.04	.11	.10	.14	.11	.54
.	.43	.17	.20	.14	.41	.40	.07	.11	.00	.13	.00	.00	.80
	.09	.07	.01	.02	.14	.10	.02	.04	.05	.04	.04	.05	.23
+	.04	.04	.04	.04	1.	1.	.04	1.	1.	.00	.00	.00	2.00
o	.19	.04	.13	.11	.10	.08	.10	.02	.01	.03	.02	.01	.31
o	.18	.30	.28	.24	.01	.21	.20	.21	.27	.10	.21	.27	.77
R	.10	.10	.19	.11	.56	.56	.02	.11	.2	.17	.89	.80	1.49
o	.15	.27	.00	.06	.20	.34	.10	.13	.17	.10	.13	.17	.61
/	.12	.08	.01	.02	.15	.09	.05	.07	.06	.06	.07	.06	.27
χ	.07	.04	.11	.05	.60	.40	.00	.00	.00	.11	1.	1.	1.60
z	.13	.28	.31	.21	.12	.07	.11	.05	.05	.20	.05	.05	.56

D = 4.22

Die größten Unterschiede zwischen beiden Sprechern bestehen in der Verwendung der ingressive Lautung, im Ersatzes des R-Lautes durch einen Reibelaut, in Vokalrundungen, im Gebrauch des uvularen R-Lautes, der Silbigkeit und der Entstimmlichung. Am wenigsten differiert die Distribution bei der Nasalisierung, den Pausen und Staupausen. Eine mittlere Distanz nimmt die Kombinatorik des schwachtonigen E-Lautes, des Glottiseinsatzes, der

Lautlänge, der vokalischen und konstriktiven (stimmhafter uvularer Reibelaut) Ersatzlautung des /R/ und der Aspiration ein. [ʀ ° χ] und [ʀ] werden jeweils nur von einem Sprecher gesprochen.

Die silbige Aussprache wird vom Sprecher weniger symmetrisch kombiniert als von der Sprecherin; die Assoziativität ist dagegen bei der Sprecherin höher. Der Sprecher kombiniert die Silbigkeit vor allem attraktiv, die Sprecherin stärker aggressiv als attraktiv. Demgemäß ist die interne Aggressivität der silbigen Verbindungen bei der Sprecherin höher, die interne Attraktivität niedriger als beim Sprecher.

Entstimmlichte Laute werden vom Sprecher mehr attraktiv als aggressiv kombiniert; konträr zur Verwendung durch die Sprecherin. Die Symmetrie ist bei beiden gleich hoch (anders als die qualitative und quantitative Symmetrie, die später noch dargestellt wird). Schwachtoniges /E/ ist in Symmetrie und Assoziativität zwischen den Sprechern um nur .1 bzw. .13 unterschiedlich; es wird von der Sprecherin eher aggressiv kombiniert.

Anders ist die Relation beim Glottisschlag: Die Symmetrie unterscheidet sich um .11, die Assoziativität um .21; die Sprecherin verwendet den Stimmeinsatz eher aggressiv-kombinatorisch, der Sprecher gleichermaßen attraktiv als auch aggressiv.

Von beiden Sprechern wird die Behauchung symmetrisch kombiniert (vgl. jedoch die qualitative Symmetrie). Der Grad der Assoziativität ist um .31 unterschiedlich. Die Sprecherin verwendet vorwiegend attraktive Kopplungen, aggressive Kombinationen treten bei beiden Sprechern gleich häufig auf.

Die Bewertung der Symmetrie bezog sich bisher auf die Relation eines Symbols i zum Symbol j, d.h. wenn die beiden Paare ij und ji im Material auftreten. Wird ein Symbol i danach klassifiziert, wieviel gemeinsame Elemente in Ai und iB existieren, erhält man eine qualitativ-symmetrische Einordnung (ALTMANN, LEHFELDT 1980, 266ff.). Die quantitativ-symmetrische Klassifikation geht aus den Differenzen der Vorgänger- und Nachfolgermenge zum Symbol i hervor. Die Artikulationstendenzen der beiden Sprecher wiesen bestimmte symmetrische Kombinationen auf, die aber nichts aussagen über die generelle Behauchung vieler Laute (d.h. etwa ein habituellder Trend) oder über Konzentrationen

auf bestimmte Laute, wie es tendentiell in der deutschen Standardaussprache der Fall ist. Die individuelle Behandlung der Kombinatorik kann durch die Zahl der Vorgänger- und Nachfolgersymbole und der in den Mengen enthaltenen Lauttypen bestimmt werden.

Qualitative und quantitative Symmetrie wurden nach folgenden Algorithmen bestimmt:

AS (qualitativ-asymmetrisch):

wenn $[Ai \cap iB] < E(Y)$ und $P(Y \leq [Ai \cap iB]) \leq \alpha$

SS (qualitativ-semisymmetrisch):

wenn $[Ai \cap iB] < E(Y)$ und $P(Y \leq [Ai \cap iB]) > \alpha$ oder

wenn $[Ai \cap iB] > E(Y)$ und $P(Y \geq [Ai \cap iB]) > \alpha$

S (qualitativ-symmetrisch)

wenn $[Ai \cap iB] > E(Y)$ und $P(Y \geq [Ai \cap iB]) \leq \alpha$

$E(Y) = (([iB] - G)([Ai] - G)) / (K - 1)$

$P(Y \leq [Ai \cap iB]) = \sum_{y_{\min}}^{[Ai \cap iB]} \left[\frac{([Ai] - G)(K - 1 - [Ai] + G)}{([iB] - y)} \right] / \binom{K-1}{[iB] - G}$

$y_{\min} = [Ai] + [iB]$ wenn $[Ai] + [iB] > K$
 $y_{\min} = G$ wenn $[Ai] + [iB] \leq K$

$P(Y > [Ai \cap iB]) = \frac{\min([Ai], [iB])}{y = [Ai \cap iB]} \left[\frac{([Ai] - G)(K - 1 - [Ai] + G)}{([iB] - y)} \right] / \binom{K-1}{[iB] - G}$

Die Klassifikation in die Kategorien quantitativ-asymmetrisch (QAS) oder quantitativ-symmetrisch (QS) erfolgte anhand der kritischen Werte für u_α (Tabelle in ALTMANN, LEHFELDT 1980, 276).

Tabelle 16: Klassifikationen der qualitativen und quantitativen Symmetrie

	SIE		ER	
n	AS	QAS	SS	QS
o	SS	QAS	S	QAS
o	SS	QS	-	-
~	SS	QS	SS	QS
!	SS	QAS	SS	QAS
!	S	QS	SS	QS
.	SS	QAS	SS	QAS
+	-	-	SS	QS
ø	SS	QS	S	QS
R	-	-	SS	QS
B	SS	QS	SS	QS
D	SS	QS	SS	QAS
χ	SS	QS	-	-
?	SS	QS	S	QS
	SS	QS	SS	QS
/	S	QS	SS	QS

lauten und Vokalen ist zwar "ungewöhnlich", kennzeichnet in seltenen Fällen jedoch ein koartikulatorisches "tonloses Sprechen", vor allem bei Vokalen ohne Phonation. Bei offenen Vokalen und stimmlosen Verschlusslauten des Sprechers sind diese Kombinationen nicht aufgetreten.

Vokale werden bevorzugt von der Sprecherin nasaliert und gerundet (Sprecher: aktuell bzw. virtuell). Stimmlose Klusile sind statistisch aktuell "gerundet oder nasaliert" notiert worden, was wiederum unüblich ist und als Versuch interpretiert werden kann, eine typische, koartikulatorische Überformung der Vokalnachfolger zu kennzeichnen.

Die Sprecherin realisiert im Gegensatz zum Sprecher auch stimmhafte Klusile und stimmlose Reibelauten silbig ($\leftrightarrow$ unzulässig). Erwartbar (virtuell) sind sogar silbig verwendete Glottiseinsätze und stimmlose Klusile (Sprecher).

Längenunterschiede treten bei stimmlosen Konstruktiven (SIE: bevorzugt, ER: aktuell), Nasale (aktuell $\leftrightarrow$ marginal), R-Allophonen (aktuell $\leftrightarrow$ unzulässig) und Glottiseinsätzen (marginal $\leftrightarrow$ unzulässig) auf; die Sprecherin spricht in allen Fällen etwas "ge-dehnter".

Ingressive Lautbildungen treten bei stimmlosen Reibelauten und Nasalen auf (nur Sprecher).

Das schwachtonige /E/ wird im Unterschied zur Sprecherin vom Sprecher vor offenen Vokalen, Glottiseinsätzen (jeweils aktuell $\leftrightarrow$ marginal) häufiger, seltener vor stimmhaften Reibelauten (aktuell $\leftrightarrow$ bevorzugt), Nasalen (marginal $\leftrightarrow$ aktuell) und Pausen (aktuell $\leftrightarrow$ bevorzugt) gesprochen.

R-Allophone spricht ER bevorzugt vor Pausen (Dialekteinfluß), aktuell ($\leftrightarrow$ marginal) vor stimmlosen Konsonanten, aber nur marginal ($\leftrightarrow$ aktuell) vor geschlossenen Vokalen.

Geschlossene Vokale werden nur marginal mit [ʔ] begonnen (ER), ($\leftrightarrow$ aktuell); stimmlose Konsonanten haben aktuell ($\leftrightarrow$ marginal) "glottale Stauungen" vor dem eigentlichen Lautbeginn, sie wären auch vor stimmlosen Klusilen zu erwarten (virtuell).

Pausen werden von beiden Sprechern ähnlich plazierte; der Sprecher bevorzugt ($\leftrightarrow$ aktuell) sie vor stimmlosen Klusilen und verwendet sie aktuell ($\leftrightarrow$ bevorzugt) vor offenen Vokalen.

Tabelle 18: Klassifikationen der Vorgängersymbole (1.: SIE, 2.: ER)

vor steht +	•	°	•	:	+	ə	RBVX	?	/
ofVo	A A	P P	A V	A P	A A	V V	M M	P P	M M
geVo	A A	A A	A V	A A	A A	V V	M M	P P	A A
shKl	M V	M M	A V	P A	A P	V V	P P	A A	M M
slKl	A A	M M	V V	A P	P A	V V	A P	A V	A A
shKo	A V	A A	A P	A A	P P	V V	A A	M M	M M
slKo	A P	P A	A V	A A	M A	V A	M M	A P	P P
Nas	P I	A A	A V	A A	A A	V A	A A	A A	A P
RBVX	P A	A A	V V	A A	P A	V V	A I	A A	A P
?	I V	M V	A V	A V	A A	V V	I V	P A	P A
/	A A	A P	A V	A I	M M	V V	M M	P P	M M

ofVo: offene Vokale geVo: geschlossene Vokale shKl: stimmhafte Klusile slKl: stimmlose Klusile
shKo: stimmhafte Konstruktive slKo: stimmlose Konstruktive Nas: Nasale

Der Sprecher spricht keine Nasale vor einem behauchten Laut ($\Leftrightarrow$ bevorzugt), beginnt nach einer Pause oder Stauung nicht mit silbigen Realisationen ($\Leftrightarrow$ aktuell) und kombiniert keine R-Allophone miteinander ($\Leftrightarrow$ aktuell). Er spricht bevorzugt stimmhafte Reibelaute vor nasalierten Vokalen ($\Leftrightarrow$ aktuell), stimmlose Verschlußlaute vor silbiger Aussprache ($\Leftrightarrow$ aktuell), stimmhafte Klusile vor langen Lauten, schwachtoniges /E/ nach stimmlosen Klusilen ($\Leftrightarrow$ aktuell; der einzige Unterschied im Gebrauch von [ə] zwischen den Sprechern), R-Allophone nach stimmlosen Verschlußlauten ($\Leftrightarrow$ aktuell), Glottiseinsätze nach stimmlosen Verschlußlauten und beginnt nach Pausen häufig mit Nasalen oder R-Allophenen ($\Leftrightarrow$ aktuell).

Die Sprecherin kombiniert bevorzugt behauchte Laute mit stimmlosen Klusilen und R-Allophenen ($\Leftrightarrow$ aktuell), Entstimmlichungen nach stimmlosen Reibelauten und Pausen ($\Leftrightarrow$ aktuell), silbige Aussprache nach stimmhaften Klusilen, Glottiseinsätze nach Pausen und mit weiteren Glottisschlägen (jeweils $\Leftrightarrow$ aktuell).

3. ZUSAMMENFASSUNG

Die statistische Analyse von auditiven Segmentierungen, der Häufigkeit und Kombination sowie der Häufigkeiten von Kombinationen hat gezeigt, daß mit dem IPA-System Unterschiede zwischen Sprachen, Formstufen und individuellen Artikulationsstilen dokumentiert werden können. Die Transkriptionsregeln sind weit genug gefaßt, um Höreindrücke von lautsprachlichen Sequenzen objektivieren zu können. Die Segmentation von Lauten ist typisch für die phänomenale Verarbeitung des akustischen Signals; auditive Zeitquantelungen sind Ausdruck der apperzeptiven Zeitstruktur. Phonetische Erwartung und fachliches Vorwissen beeinflussen nicht übermäßig stark die Transkription.

Die Umgangssprache unterscheidet sich von der schriftsprachlichen Kombinatorik und Lauthäufigkeit durch einen Anstieg der Kombination stimmhafter Konsonanten mit anderen Lautgruppen und durch eine Abnahme von Kombinationen mit stimmlosen Konsonanten, vor allem mit stimmlosen Verschlußlauten. Liquide werden häufig reduziert und verschliffen ausgesprochen; Nasale und Vokale sind am wenigsten beeinflusst.

Sprechertypisch und individuell war der Gebrauch der Aspiration, der Entstimmlichung, der Glottiseinsätze, der Nasalierung und der Reduktionen unbetonter Endsilben. Pausengestaltung und die Silbigkeit von Lauten unterschieden sich nicht sprecherspezifisch.

Fußnoten:

- 1 Zwei Unterschiede zum IPA-Standard sind zu beachten:
[o] steht für [ɔ]; [ɐ] für [ʊ].
- 2 Herrn Wolfgang Brudler danke ich herzlich für die Auszählung und Tabellierung der Daten.

Anhang

Programm Distributionsanalyse (TI 59, PC 100 C, ML-Modul)

Labels: 001 E'
 013 A'
 019 A
 024 B
 029 C
 034 C'
 039 E

User defined labels and user instruction:

1 Distribution			2	
ii?		K (1x)		Entry
Ai	iB	A/AB		Code Sy

Use order E' - C' (first)

(A' if ii)

A+ B+ C+ E (starts program)

Entry	Store K	(SD in Reg 03)	(SD)
000 LBL	033 LBL	067 =	104 3
001 E'	034 C'	068 =	105 6
002 PGM	035 STO	069 (	106 1
003 01	036 07	070 2	107 6
004 SBR	037 R/S	071 x	108 OP
005 CLR		072 RCL	109 04
006 0	Store code/	073 02	110 RCL
007 STO	run pgm	074)	111 03
008 00	038 LBL	075 =	112 OP
009 OP	039 E	076 STO	113 06
010 00	040 OP	077 03	114 ADV
011 R/S	041 02		
	042 OP	(VE)	(At)
	043 05	078 4	115 1
Store G		079 2	116 3
012 LBL		080 1	117 3
013 A'	(Ai)	081 7	118 7
014 1	044 1	082 OP	119 OP
015 STO	045 3	083 04	120 04
016 05	046 7	084 RCL	121 RCL
017 R/S	047 4	085 02	122 00
	048 OP	086 +	123 ÷
Store Ai	049 04	087 RCL	124 RCL
018 LBL	050 RCL	088 03	125 07
019 A	051 00	089 =	126 =
020 STO	052 OP	090 STO	127 OP
021 00	053 06	091 04	128 06
022 R/S		092 OP	
	(iB)	093 06	(Ag)
Store iB	054 7		129 1
023 LBL	055 4		130 3
024 B	056 1	(DU)	
025 STO	057 4	094 1	131 2
026 01	058 OP	095 6	132 2
027 R/S	059 04	096 4	133 OP
	060 RCL	097 1	134 04
Store A/AB	061 00	098 OP	135 RCL
028 LBL	062 +	099 04	136 01
029 C	063 RCL	100 RCL	137 ÷
030 STO	064 01	101 02	138 RCL
031 02	065 OP	102 OP	139 07
032 R/S	066 06	103 06	140 =
			141 OP
			142 06

(As)	195 OP	(S'*)	299 5
143 1	196 06	i	300 2
144 3	197 RCL	247 3	301 4
145 3	198 00	248 6	302 OP
146 6	199 ÷	249 2	303 04
147 OP	200 RCL	250 4	304 1
148 04	201 04	251 5	305 -
149 RCL	202 =	252 1	306 (
150 04	203 OP	253 OP	307 RCL
151 ÷	204 06	254 04	308 02
152 RCL		255 RCL	309 ÷
153 07	(Ag _i)	256 02	310 RCL
154 =	205 1	257 x	311 04
155 OP	206 3	258 2	312)
156 06	207 2	259 =	313 =
	208 2	260 ÷	314 OP
(As*)	209 2	261 (	315 06
157 1	210 4	262 RCL	
158 3	211 OP	263 00	(S _i '*)
159 3	212 04	264 +	316 3
160 6	213 RCL	265 RCL	317 6
161 5	214 01	266 01	318 6
162 1	215 ÷	267)	319 5
163 OP	216 RCL	268 =	320 2
164 04	217 04	269 STO	321 4
165 RCL	218 =	270 08	322 5
166 00	219 OP	271 OP	323 1
167 +	220 06	272 06	324 OP
168 RCL			325 04
169 01	(S)	(S')	326 1
170 =	221 3	273 3	327 -
171 -	222 6	274 6	328 RCL
172 RCL	223 OP	275 6	329 08
173 05	224 04	276 5	330 =
174 =	225 RCL	277 OP	331 OP
175 ÷	226 02	278 04	332 06
176 (	227 ÷	279 RCL	
177 (	228 RCL	280 04	(G)
178 2	229 07	281 ÷	333 2
179 x	230 =	282 RCL	334 2
180 RCL	231 OP	283 07	335 OP
181 07	232 06	284 =	336 04
182)		285 -	337 RCL
183 -	(S _i)	286 (	338 05
184 1	233 3	287 RCL	339 OP
185)	234 6	288 02	340 06
186 =	235 2	289 ÷	341 ADV
187 OP	236 4	290 RCL	342 ADV
188 06	237 OP	291 07	343 ADV
	238 04	292)	344 ADV
(At _i)	239 RCL	293 =	345 GTO
189 1	240 02	294 OP	346 E'
190 3	241 ÷	295 06	
191 3	242 RCL		Register Contents
192 7	243 04	(S _i ') i	01 Ai 05 G
193 2	244 =	296 3	02 iB 06
194 4	245 OP	297 6	03 DU 07 K
	246 06	298 6	04 VE 08 used

Programm Klassifikation quantitative und qualitative Symmetrie

(NB: im Anschluß an die Distributionsanalyse kann das Programm eingelesen werden und mit Lbl E gestartet werden. Die Rechenzeit kann stark variieren, je nach Größe der Zahlen, für die eine Kombinatorik berechnet werden muß.)

000	LBL	050	GE	100	RCL	150	STO
001	E	051	E'	101	08	151	13
002	RCL	052	1	102	)	152	IFF
003	02	053	SUM	103	x	153	02
004	STO	054	15	104	(	154	B'
005	12	055	3	105	RCL	155	RCL
006	RCL	056	0	106	10	156	10
007	00	057	GE	107	-	157	+
008	STO	058	E'	108	RCL	158	RCL
009	10	059	1	109	08	159	11
010	RCL	060	SUM	110	)	160	=
011	01	061	15	111	=	161	x ott
012	STO	062	3	112	÷	162	RCL
013	11	063	8	113	RCL	163	07
014	RCL	064	GE	114	07	164	GE
015	05	065	E'	115	=	165	SIN
016	STO	066	1	116	STO	166	x ott
017	08	067	SUM	117	09	167	-
018	RCL	068	15	118	x ott	168	(
019	07	069	4	119	RCL	169	RCL
020	x ott	070	8	120	12	170	07
021	PGM	071	GE	121	EQ	171	+
022	01	072	E'	122	TAN	172	1
023	SBR	073	1	123	GE	173	)
024	CLR	074	SUM	124	ENG	174	=
025	4	075	15	125	GTO	175	STO
026	STO	076	5	126	SUM	176	05
027	15	077	8	127	LBL	177	GTO
028	7	078	GE	128	ENG	178	COS
029	GE	079	E'	129	STF	179	LBL
030	E'	080	1	130	02	180	SIN
031	1	081	SUM	131	LBL	181	RCL
032	SUM	082	15	132	SUM	182	08
033	15	083	7	133	RCL	183	STO
034	1	084	0	134	07	184	05
035	1	085	GE	135	PGM	185	LBL
036	GE	086	E'	136	16	186	COS
037	E'	087	1	137	A	187	RCL
038	1	088	SUM	138	RCL	188	12
039	SUM	089	15	139	11	189	-
040	15	090	LBL	140	-	190	RCL
041	1	091	E'	141	RCL	191	05
042	6	092	1	142	08	192	+
043	GE	093	INV	143	=	193	1
044	E'	094	SUM	144	PGM	194	=
045	1	095	07	145	16	195	STO
046	SUM	096	(	146	B	196	06
047	15	097	RCL	147	PGM	197	GTO
048	2	098	11	148	16	198	D
049	3	099	-	149	E	199	LBL

200	B'	256	STO	312	OP	369	3
201	RCL	257	14	313	03	370	6
202	12	258	RCL	314	GTO	371	OP
203	STO	259	07	315	C	372	03
204	05	260	-	316	LBL	373	LBL
205	RCL	261	RCL	317	NOP	374	PRT
206	10	262	10	318	IFF	375	OP
207	x ott	263	+	319	02	376	05
208	RCL	264	RCL	320	GRD	377	ADV
209	11	265	08	321	1	378	OP
210	GE	266	=	322	3	379	00
211	FIX	267	PGM	323	3	400	CLR
212	STO	268	16	324	6	401	R/S
213	06	269	A	325	OP	402	LBL
214	GTO	270	RCL	326	03	403	DEG
215	OP	271	11	327	GTO	404	3
216	LBL	272	-	328	C	405	4
217	FIX	273	RCL	329	LBL	406	1
218	x ott	274	05	330	GRD	407	3
219	STO	275	=	331	3	408	3
220	06	276	PGM	332	6	409	6
221	LBL	277	B	333	OP	410	OP
222	OP	278	PGM	334	03	411	03
223	RCL	279	16	335	LBL	412	GTO
224	06	280	E	336	C	413	PRT
225	-	281	x	337	1		
226	RCL	282	RCL	338	SUM		
227	05	283	14	339	07		
228	+	284	=	341	OP		
229	1	285	÷	342	05		
230	=	286	RCL	343	PGM		
231	STO	287	13	344	01		
232	06	288	=	345	SBR		
233	LBL	289	SUM	346	CLR		
234	D	290	00	347	0		
235	RCL	291	1	348	STO		
236	10	292	SUM	349	00		
237	-	293	05	350	STO		
238	RCL	294	DSZ	351	08		
239	08	295	06	352	INV		
240	=	296	D	353	STF		
241	PGM	297	.	354	02		
242	16	298	0	355	RCL		
243	A	299	5	356	15		
244	RCK	300	x ott	357	x ott		
245	05	301	RCL	358	RCL		
246	-	302	00	359	10		
247	RCL	303	INV	360	-		
248	08	304	GE	361	RCL		
249	=	305	NOP	362	11		
250	PGM	306	LBL	363	=		
251	16	307	TAN	364	IxI		
252	B	308	3	365	GE		
253	PGM	309	6	366	DEG		
254	16	310	3	367	3		
255	E	311	6	368	4		

Programm Klassifikation nach distributioneller Frequenz (ML-Modul)

(NB: Bei sehr kleinem P kann der Rechner blinken ohne das Programm zu unterbrechen; vor Neueingabe CLR drücken. Lbl E: für "O" vor-gesehen; bei mehreren gleichen Werten kann die Zahl aktuell in das Rechenprogramm eingesetzt werden - manuell GTO o22, LRN, Zahl ändern, LRN. Σ_{ij} bei Ersteingabe mittels B speichern, sonst nur C.
Zeitbedarf: ca. 11 Sek.; bei Matrix 50x50: 1 Spalte 9, 1 Minuten)

Labels:

o19 E
o32 C
198 A'
2o8 A
212 E'
218 B
223 C'

User defined labels and user instruction:

1 Distributional Frequency Top17 2				
Symbol		n		Entry
$\Sigma_i.$	$\Sigma.j$	Σ_{ij}		$\Sigma_{ij}=0?$

Use order C' → A → A' → C₁ (or C₁ → B₁
C₂ C₂ → B₂
C₃ C₃ → B₃
etc. etc.)

(SBR categories) start

ooo RC* o31 LBL
oo1 69 o32 C
oo2 OP o33 STO
oo3 o4 o34 56
oo4 RCL o35 CP
oo5 57 o36 INV
oo6 - o37 EQ
oo7 4 o38 oo
oo8 = o39 45
oo9 OP o4o 1
o1o o6 o41 SUM
o11 1 o42 69
o12 SUM o43 STF
o13 57 o44 o1
o14 o o45 RC*
o15 STO o46 57
o16 69 o47 x
o17 RTN o48 RCL
o49 58
Store counter o5o =
o18 LBL o51 ÷
o19 E o52 RCL
o2o STO o53 59
o21 55 o54 =
o22 o o55 STO
o23 SBR o56 67
o24 oo o57 NOP/PRT
o25 33 o58 x↔t
o26 DSZ o59 2
o27 55 o6o 9
o28 oo o61 GE
o29 22 o62 o1
o3o R/S o63 41

o64 (o98 RCL
o65 RCL o99 59
o66 56 1oo -
o67 - 1o1 1
o68 x↔t 1o2)
o69) 1o3)
o7o ÷ 1o4)
o71 (1o5 $\sqrt{x}$
o72 (1o6)
o73 (1o7 =
o74 RCL 1o8 NOP/PRT
o75 59 1o9 x↔t
o76 - 11o RCL
o77 RCL 111 66
o78 58 112 INV
o79) 113 GE
o8o x 114 o1
o81 (115 27
o82 RCL 116 +/-
o83 59 117 GE
o84 - 118 o1
o85 RC* 119 34
o86 57 12o 6
o87) 121 2
o88 x 122 SUM
o89 RCL 123 69
o9o 67 124 GTO
o91) 125 oo
o92 ÷ 126 oo
o93 (127 6
o94 RCL 128 o
o95 59 129 SUM
o96 x 13o 69
o97 (131 GTO

132 oo 189 RCL
133 oo 19o 56
134 6 191 GE
135 4 192 o1
136 SUM 193 27
137 69 194 GTO
138 GTO 195 o1
139 oo 196 34
14o oo
141 IFF Store Code
142 o1 197 LBL
143 o1 198 A'
144 6o 199 OP
145 RCL 2oo oo
146 56 2o1 OP
147 PGM 2o2 o1
148 16 2o3 ADV
149 A 2o4 OP
15o PGM 2o5 o5
151 16 2o6 R/S
152 C
153 RCL Store $\Sigma_i.$
154 67 2o7 LBL
155 y* 2o8 A
156 RCL 2o9 STO
157 56 21o 58
158 = 211 LBL
159 x 212 E'
16o RCL 213 5
161 67 214 STO
162 +/- 215 57
163 INV 216 R/S
164 LNXX
165 = Store $\Sigma.j$
166 IFF 217 LBL
167 o1 218 B
168 o1 219 ST*
169 75 22o 57
17o ÷ 221 R/S
171 RCL
172 o4 Store n
173 = 222 LBL
174 NOP/PRT 223 C'
175 x↔t 224 STO
176 INV 225 59
177 STF 226 R/S
178 o1
179 .
18o o
181 5
182 INV
183 GE
184 o1
185 2o
186 RCL
187 67
188 x↔t

Register Contents:

oo - o4 used (PGM 16)
o4 n!
o5 - 54 Σ_j
55 counter
56 Σ_{ij}
57 counter
58 Σ_i
59 n
6o-65 code categories
66 1.645
67 E
68 z
69 counter for 6o-65

flag o1
x↔t

Store following codes
and date up on card side 2:
6o 33oooo
61 17353532
62 13oooo
63 42oooo
64 3ooooo
65 24oooo
66 1.645

LITERATURVERZEICHNIS

ALTMANN, G.

- 1971 Introduction to Quantitative Phonology. Habilitations-schrift Bochum.

ALTMANN, G., LEHFELDT, W.

- 1972 Typologie der phonologischen Distributionsprofile. Beiträge zur Linguistik und Informationsverarbeitung 22, 8-32.
- 1980 Einführung in die Quantitative Phonologie. Bochum: Brockmeyer.

ANDREEV, N.D. (Hrsg.)

- 1965 Statistiko-kombinatornoe modelirovanie jazykov. Moskva-Leningrad: Nauka.

DENES, P.D.

- 1964 On the statistics of spoken English. Zeitschrift für Phonetik, Sprachwissenschaft und Kommunikationsforschung 17, 51-71.

KARLGREN, H.

- 1974³ Statistical methods in Phonetics. In MALMBERG, B. (Ed.): Manual of Phonetics. Amsterdam, London, New York, 129-154.

LADEFOGED, P.

- 1975 A Course in Phonetics. New York.

LEHFELDT, W.

- 1973 Distributionelle Phonemähnlichkeit. Phonetica 27, 82-99.
- 1975 Die Verteilung der Phonemanzahl in den natürlichen Sprachen. Phonetica 31, 274-287.
- 1976 Zur Messung der phonetischen Lautdifferenz. Eine begriffskritische Untersuchung. Glottometrika 1, 26-45.

LEHFELDT, W., ALTMANN, G.

- 1975 Begriffskritische Untersuchungen zur Sprachtypologie. Linguistics 144, 49-78.

LINDNER, G.

- 1975 Der Sprechbewegungsablauf. Eine phonetische Studie des Deutschen. Berlin: Akademie-Verlag.
- 1980 Lautfolgestrukturen im Deutschen. Zeitschrift für Phonetik, Sprachwissenschaft und Kommunikationsforschung 33, 468-477.

OLLER, D.K., EILERS, R.E.

- 1975 Phonetic expectation and transcription validity. Phonetica 31, 288-304.

ORTMANN, W.D.

- 1975 Beispielwörter für deutsche Ausspracheübungen. München.

POYATOS, F.

- 1975 Cross-Cultural Study of Paralinguistic "Alternants" in Face-to-Face-Interaction. In KENDON, A., HARRIS, R.M., KEY, M.R. (Eds.): Organization of Behaviour in Face-to-Face-Interaction, 285-313.

RICHTER, B., RICHTER, H.

- 1981 Zur Abbildungstreue von Transkriptionen. In LANGE-SEIDL, A. (Hrsg.): Zeichenkonstitution II, 110-119.

SEGAL, M.

- 1972 Osnovy fonologičeskoj statistiki. Moskva-Leningrad: Nauka.

TIFFANY, W.R., CARRELL, J.

- 1977² Phonetics. New York.

WOTHKE, K.

- 1980 Computerunterstützte Phonemanalyse. In HALLER, J. (Hrsg.): Maschinelle Sprachverarbeitung, 36-58.

WINKLER, P.

- 1979 Notationen des Sprechausdrucks. Zeitschrift für Semiotik 1, 211-224.
- 1981 Anwendungen phonetischer Methoden für die Analyse von Face-to-Face-Situationen. In WINKLER, P. (Hrsg.): Methoden der Analyse von Face-to-Face-Situationen. Stuttgart: Metzler, 9-46.
- 1982 Notation of Paraphonetic Features. In HESS-LÜTTICH, E.W.B. (Ed.): Multimedial Communication 1, Tübingen: Narr.

GRAPHEME UND GRAPHEMKOMBINATORIK DER RUSSISCHEN FACHSPRACHE

EINE PHONOSTATISTISCHE UNTERSUCHUNG

J. Dietze, Halle (Saale)

Die Phonostatistik beschäftigt sich mit dem frequentativen Verhalten von Phonem (Lauten), Phonemen sowie distinktiven Phonemmerkmalen, um entsprechende quantitative Regularitäten zu ermitteln (vgl. ROCZAWSKI 1876, 21). Von ihr zu unterscheiden ist die Phonotaktik, deren Aufgabenbereich APRESJAN (1971, 52) wie folgt definiert: "Innerhalb der Phonologie sei die Phonotaktik herausgehoben, die Wissenschaft von den Gesetzen der Phonemverbindung, welche die strukturbedingten Beschränkungen der Kombinationsmöglichkeiten der Phoneme erforscht." Für die Untersuchung phonotaktischer Fragestellungen lassen sich statistische Methoden erfolgreich einsetzen. Phoneme sind die kleinsten phonologischen Einheiten, die nicht in noch kleinere aufgespalten werden können (vgl. KOSCHMIEDER 1977, 13). Strukturell gesehen, entspricht dem Phonem auf der graphematischen Ebene das Graphem. Das Graphem ist das Bezeichnende, das Phonem das Bezeichnete. Zwischen dem Graphemsystem einer Sprache und deren Phonemsystem bestehen Beziehungen, die auch direkte Schlüsse vom Graphemsystem auf das Phonemsystem gestatten können; denn die phonologische Struktur determiniert die graphematische, wenn auch vermittelt. Es wäre deshalb korrekt, die folgende Untersuchung "graphostatistisch" zu nennen. Wegen der Ungewöhnlichkeit dieses Terminus wollen wir sie jedoch - unter Beachtung der genannten Einschränkungen - der Phonostatistik zuordnen.

Die Buchstaben des Russischen nennen wir "Grapheme". Da in den von uns ausgewerteten Texten keine Allographie vorkommen (Großbuchstaben und $\bar{e}$ wurden als Kleinbuchstaben bzw. e abgelocht), brauchen wir die begriffliche Unterscheidung von Graphemen und Graphen in unserem Falle nicht zu beachten.¹

Die erste umfassende Arbeit zur Phonostatistik des Russischen haben BELONOGOV und FROLOV (1963) veröffentlicht. Diesen Autoren hatten russische Texte der Geschäftssprache mit einem Umfang von 30 000 Wörtern bzw. 200 000 Buchstaben als Materialgrundlage gedient. TOPOROV (1966, 110) hält den Umfang dieser Stichprobe für zu klein - zu Unrecht, wie wir meinen. Jedenfalls wird für sie keine nachprüfbare Begründung angeführt. Ferner sind TOPOROVs eigene Untersuchungen ohne textstatistische Parameter, da sein Material erklärenden Wörterbüchern entstammt, Quellen also, in denen nichts über die Auftretenshäufigkeit der verzeichneten Einheiten zu finden ist. Um festzustellen, ob eine Stichprobe für sprachstatistische Forschungen hinreichend groß ist, kann man nach der Wahrscheinlichkeitsrechnung den relativen Fehler approximativ ermitteln. Nach FRUMKINA (1973, 282) benützt man hierfür folgenden Ausdruck:

$$\frac{z_p}{NP}$$

N = Stichprobenumfang (beispielsweise Anzahl der Buchstaben)

P = relative Häufigkeit der zu untersuchenden Einheit

z_p = Konstante, die die Größe der Wahrscheinlichkeit angibt. Wir setzen hier $z = 1.96$, damit die Konfidenz 95 % beträgt (vgl. KALININA 1973, 94).

Wenn man P als arithmetisches Mittel der Vorkommenshäufigkeit eines russischen Graphems in der relativen Form (bezogen auf die Gesamtzahl der Buchstaben) auffaßt, dann beträgt der mittlere relative Fehler 0.0248, d.h. rund 2.5 %. Bei sprachstatistischen Untersuchungen sollte der relative Fehler 30 % nicht übersteigen.

Unserer phonostatistischen Untersuchung wurden 500 russischsprachige Referate einschließlich ihrer Titel zugrundegelegt, die der Referatezeitschrift (Referativnyj žurnal) zum Bereich "wissenschaftliche und technische Information" entstammen. Dieses Textkorpus enthält 57 566 Wortformen mit insgesamt 429 257 Buchstaben, wobei hier, wie erwähnt, der Unterschied zwischen

Groß- und Kleinbuchstaben vernachlässigt wurde. Das arithmetische Mittel der Frequenz der russischen Grapheme beträgt in unserem Falle 13 414.28. Diese Zahl wird in die obige Formel eingesetzt, so daß sich der mittlere relative Fehler auf 0.0169, d.h. rund 1.7 %, beläuft. Die Größe der von uns gewählten Stichprobe dürfte in Anbetracht des Untersuchungsgegenstandes statistisch repräsentativ sein; denn die kleinste Frequenz von 156 (für Ъ) weist nur einen relativen Fehler von 0.157, d.h. 15.7 % auf. Zum Vergleich sei auf die Arbeit von GVOZDOVIČ (1976c) hingewiesen. Dieser Autor hat deutschsprachige Texte mit 50 000 Buchstaben bezüglich der Distribution untersucht und gibt eine mögliche Fehlerquote von 5 % an.

GVOZDOVIČ (GWOSDOWITSCH 1976a) hat in gleicher Weise wie TOPOROV neben ausländischen Ortsnamen auch Abkürzungen in seine graphemkombinatorischen Untersuchungen einbezogen. Soweit es sich bei den Abkürzungen um Kurzwörter handelt, die nach dem Vorbild der Akronyme aus den Anfangssilben der abzukürzenden Wortformen gebildet werden, ist gegen dieses Verfahren nichts einzuwenden, da bei solchen Abkürzungen das phonotaktische System nicht verletzt wird. Werden jedoch Akronyme aus den Anfangsbuchstaben der abzukürzenden Wortformen zusammengesetzt, so sollten sie bei phonostatistischen Forschungen insofern nicht berücksichtigt werden, als sie die Untersuchungsergebnisse verfälschen können. Beim Ablochen der Texte haben wir bereits gängige fachliche Akronyme aufgelöst.

In Tabelle 1 führen wir die Vorkommenshäufigkeit der einzelnen russischen Grapheme an, und zwar in absoluten sowie in Prozentzahlen. Ergänzend dazu bringen wir auch die in Prozentzahlen umgerechneten Werte von BELONOGOV und FROLOV (1963, 287).² Die Reihenfolge der Buchstaben ergibt sich aus der absteigenden Frequenz.

Die im rückläufigen Wörterbuch des Russischen von 1974 (Obratnyj slovar') aufgeführten phonostatistischen Angaben konnten wir leider nicht zum Vergleich heranziehen, weil sie nur den Auslaut betreffen. Gleiches gilt auch für die noch zu behandelnden russischen Graphemkombinationen.

Den phonetischen Gegensatz "hart ~ palatalisiert" der in

Tabelle 1: Absolute und relative Häufigkeit der russischen Grapheme in der Textstichprobe

Graphem	absolute Frequenz	relative Frequenz in %	
		DIETZE	BELONOGOV, FROLOV
О	44 172	10.29	10.47
И	42 024	9.79	7.00
Е	35 662	8.31	8.36
А	33 967	7.91	8.08
Н	29 877	6.96	7.23
Т	27 447	6.39	6.25
С	26 034	6.06	4.66
Р	22 279	5.19	5.84
В	17 586	4.10	5.69
Л	14 613	3.40	2.48
К	14 189	3.31	2.64
М	13 890	3.24	3.37
П	12 736	2.97	3.71
Д	11 079	2.58	3.88
Я	9 893	2.30	2.49
Ы	8 632	2.01	1.79
У	8 413	1.96	2.02
Э	7 000	1.63	1.47
Б	6 464	1.51	1.70
Ч	6 005	1.40	1.80
Х	5 390	1.26	1.32
Й	4 852	1.13	1.58
Г	4 716	1.10	1.11
Ц	4 491	1.05	0.29
Ь	4 389	1.02	1.68
Ф	3 912	0.91	0.17
Ю	2 904	0.68	0.63
Ж	2 537	0.59	0.96
Щ	1 670	0.39	0.35
Ш	1 224	0.29	0.50
Э	1 054	0.25	0.17
Ъ	156	0.04	0.03

Frage kommenden russischen Konsonanten haben wir statistisch ebenfalls auf maschinellm Wege approximativ zu erfassen versucht.³ Als palatalisiert gelten hierbei die unten angeführten Konsonanten in folgender Position: Konsonant + е, и, ъ, ю, я. Harte Konsonanten stehen in folgenden Positionen: 1. Konsonant + а, о, у, ь, э. 2. Konsonant im absoluten Auslaut. 3. Konsonant + Konsonant(en). Bei Konsonantenverbindungen konnten wir mögliche Palatalisierungen, die im Russischen meist regressiv wirken und Assimilationen, die die statistischen Werte verschieben, nicht berücksichtigen. Die von uns beigebrachten Prozentangaben beziehen sich jeweils auf die Gesamtfrequenz des betreffenden Konsonantengraphems.

Tabelle 2:

Graphem	Palatalisiert absolute Frequenz / in %		Hart absolute Frequenz / in %	
б	1243	19.2	5221	80.8
в	2655	15.0	14931	85.0
г	596	12.6	4120	87.4
д	3577	32.3	7502	67.7
э	684	9.8	6316	90.2
н	2354	16.6	11835	83.4
л	10175	69.6	4438	30.4
м	4534	32.6	9356	67.4
н	8471	28.4	21406	71.6
п	2384	18.7	10352	81.3
р	6782	30.4	15497	69.6
с	5227	20.1	20807	79.9
т	9534	34.7	17913	65.3
ф	1329	34.0	2583	66.0
х	199	3.7	5191	96.3

Addiert man zu den palatalisierten russischen Konsonanten (59 744) noch die Frequenz der Grapheme der palatalen Konsonan-

ten (ч, щ) so ergibt sich eine Frequenz von insgesamt 67 419. Addiert man in gleicher Weise die Frequenz der eben erfaßten harten Konsonanten (176 676) zu der Frequenz der Grapheme für positionsunabhängig harte Konsonanten (ж, ц, ш), so erhalten wir als Summe 185 928. Von den 253 347 Konsonantengraphemen bezeichnen mithin 26.6 % palatalisierte bzw. palatale Konsonanten und 73.4 % harte Konsonanten. Diese Angaben sind unter Beachtung der oben gemachten Einschränkungen lediglich als Näherungswerte zu verstehen.

Bei der nun folgenden Behandlung von Graphemverbindungen beschränken wir uns bewußt auf die Verbindungen von Konsonantengraphemen untereinander, unter Einbeziehung des Weichheitszeichens. BELONOGOV und FROLOV (1963) haben die bigraphemischen Verbindungen unter Einschluß der Vokalgrapheme untersucht, so daß deren statistische Frequenzangaben mit den unsrigen nicht vergleichbar sind. In gleicher Weise ist auch TOPOROV (1966) vorgegangen. Beide Arbeiten zeigen aber bei den bigraphemischen Konsonantenverbindungen mehr Kombinationsmöglichkeiten als in unseren Texten. Das trifft vor allem für die Angaben bei TOPOROV zu und ist damit zu erklären, daß TOPOROV Lexika ausgewertet hat, und zwar unter Einschluß von Akronymen (Bildung mit Hilfe der Initialbuchstaben). Das Besondere der Untersuchung TOPOROVs besteht in der Klassifizierung der Graphempaare nach Distributionsmerkmalen (Valenz, Symmetrie, Transitivität), wobei er der von HARARY und PAPER (1957) entwickelten Methode folgt, die auf der Relationstheorie beruht und mit deren Hilfe die "interaction" von Graphempaaren bestimmt wird.⁴ Wir müssen jedoch einwenden, daß damit von der Graphemebene aus kaum weiterführende Einsichten in höher liegende Strukturebenen wie die der Phonetik oder Phonologie oder gar der Morphologie gewonnen werden können. Wir haben uns daher entschlossen, die konsonantischen Graphempaare im Anhang lediglich in Form eines Überblicks darzubieten. Dies geschieht mittels einer Tabelle, in der auf der ersten Spalte die Konsonantengrapheme in erster Position und in der ersten Zeile die Konsonantengrapheme in zweiter Position eingezeichnet sind. Diese Tabelle gestattet es, den Sättigungsgrad der Valenz (= Anzahl der Grapheme, die

mit dem zu untersuchenden Graphem eine Verbindung eingehen) sowie die Symmetrie (= Fähigkeit eines Graphems, die 1. und/oder die 2. Position einzunehmen) der Grapheme zu bestimmen. Daraus lassen sich wiederum Distributionsklassen aufstellen.⁵ Unsere Frequenzangaben sind absolut, beim Errechnen der relativen Frequenz müßte man sich auf die Summe aller von uns gezählten konsonantischen Graphempaare beziehen (52 454), vgl. Anhang I.

Vom Typ her konnten wir 214 bigraphemische Konsonantenverbindungen von 400 (= 20²) mathematisch möglichen feststellen; das sind 53.5 %. TOPOROV (1966, 66) hat ermittelt, daß der Anteil der bigraphemischen Verbindungen von den rechnerisch möglichen unter Berücksichtigung aller Grapheme bei 75 % liegt.

Die Kombinatorik der Konsonantengrapheme läßt u.E. erst dann Einsichten auf der höherliegenden Strukturebene zu, wenn die ermittelten Graphemverbindungen typisiert werden. Wir folgen hier dem Versuch von ŠABROVA (1978), die die russischen Konsonantengrapheme zu phonetisch definierten Gruppen zusammengefaßt hat. Dieses Vorgehen hat den Vorteil, daß Palatalisierungserscheinungen und vor allem regressive Stimmassimilationen, die auf der phonetischen Ebene ablaufen, auf der Graphemebene neutralisiert werden. Für die Gruppenbildung werden artikulatorische Merkmale unter Beachtung der Klassifikation nach Geräuschkonsonanten und Sonoren herangezogen: T = Klusile (б, п; г, н; д, т); S = Spiranten (з, с; ж, ш; в, ф; х); Š = Sonore (м, н, р, л); A = Affrikaten (ч, ц, щ). Die kodierte Bezeichnung der Gruppen haben wir ebenfalls von ŠABROVA übernommen, um Vergleiche des Materials zu erleichtern. ŠABROVA bietet jedoch keine Frequenzangaben, sondern nur distributive Gruppierungen. In diesem Zusammenhang ist der Hinweis wichtig, daß alle bigraphemischen Gruppen mit zwei Ausnahmen zu dem sogenannten gemischten Typ gehören, d.h. die konsonantischen Graphempaare stehen sowohl auf der Morphemgrenze als auch innerhalb eines Morphems. Nur die Gruppen TA und AS kommen ausschließlich auf der Morphemgrenze vor.

Bei der maschinellen Aufbereitung unserer Texte haben wir jeweils den Wortlaut gesondert gekennzeichnet, so daß wir bei

den einzelnen Gruppen der Graphempaare die im Anlaut stehenden Verbindungen angegeben werden. Hervorheben möchten wir diejenigen konsonantischen Graphempaare, die allein im Anlaut vorkamen:⁶ сг-, сд-, сж-, сф-, тц-, фл-, шв-, шп-, шр-. Zur Illustration unserer obigen Behauptung, daß Akronyme des von uns ausgeschlossenen Bildungstyps in der Regel die graphisch reflektierten phonetischen Gesetzmäßigkeiten durchbrechen, seien die entsprechenden Abkürzungen hier angeführt, da sie zum Anlaut gehören: км, мм, р(ефератный) ж(урнал), ц(ентральный) н(аталог), Ц(ентральный) Н(аучно-) И(сследовательский) И(нститут) П(атентной) И(нформации).

Von den 16 mathematisch möglichen Gruppen der bigraphemischen Konsonantenverbindungen sind in unseren Texten 15 realisiert (93.8 %). Wir geben die absolute Frequenz der einzelnen Graphempaare, die zu der Gruppe gehören, an, darauf folgen die relative Frequenz in Prozenten (bezogen auf die Gesamtzahl der konsonantischen Graphempaare 52 454) und abschließend die Anzahl der anlautenden Graphempaare (absolute Frequenz).

Tabelle 3:

Graphemgruppe	absolute Gesamtfrequenz / in %		Anlaut
TT	2125	4.1	13
TS	3981	7.6	220
TR	11820	22.5	5083
TA	419	0.8	4
ST	9474	18.1	997
SS	2728	5.2	1055
SR	6559	12.5	1667
SA	108	0.2	26
RT	3097	5.9	-
RS	3000	5.7	-
RR	7128	13.6	104
RA	255	0.5	-
AT	420	0.8	309
AS	43	0.1	-
AR	1194	2.3	16

An dreigraphemischen Konsonantenverbindungen haben wir insgesamt 4 438 in unseren Texten gezählt. Vom Typ her wären rechnerisch $20^3 = 8\ 000$ Verbindungen möglich, 142 kommen jedoch nur vor (1.8 %), wobei die ermittelten folgenden Abkürzungen unberücksichtigt blieben: б(иблютечно-) б(иблюграфическая) н(лассификация), в(енгерская) н(ародная) р(еспублика), ГДР, Н(ародная) р(еспублика) Б(олгария), П(ольская) н(ародная) р(еспублика), р(аскладочно-) п(одборочная) м(ашина), с(четно-) п(ерфорационная) м(ашина), С(оциалистическая) р(еспублика) Р(умыния), С(оветская) С(оциалистическая) р(еспублика), Ф(едеративная) р(еспублика) Г(ермания), (электронная) ц(ифровая) в(ычислительная) м(ашина). Demgegenüber finden sich folgende dreigraphemische Konsonantenverbindungen, die keine Abkürzungen darstellen, nur im Anlaut: вэв-, вэр-, вкл-, впл-, всл-, вст-, эдр-, мгн-, сбл-, сгл-, сгр-.⁶

Von 64 (= 4^3) mathematisch möglichen Gruppen dreigraphemischer Konsonantenverbindungen finden sich 28 (43.8 %) in den Texten. Die Prozentangaben bei der Frequenz beziehen sich auf die Gesamtzahl aller dreigraphemischen Konsonantenverbindungen (4 438).

Die Summe aller viergraphemischen Konsonantenverbindungen ist 795. Vom Typ her wären rechnerisch $20^4 = 160\ 000$ Verbindungen möglich, 34 kommen jedoch nur vor (0.02 %); dies ohne Berücksichtigung der Abkürzungen СССР und ЧССР. Echte, ausschließlich im Anlaut stehende Verbindungen sind вэгл- sowie вскр-.⁶

Von 256 (= 4^4) mathematisch möglichen Gruppen viergraphemischer Konsonantenverbindungen finden sich 15 (5.9 %). Die Prozentangaben bei der Frequenz beziehen sich auf die Gesamtzahl aller viergraphemischen Konsonantenverbindungen (795).

ТОПОРОВ (1966, 126) führt noch eine Reihe von Verbindungen mit 5 bzw. 6 Konsonantengraphemen an, es handelt sich dabei aber mit folgenden Ausnahmen immer um Fremdwörter: рхвкл (сверхвключение), рхвкр (сверхвкрадчивый); dieser Typ RS/STR liegt auf der Morphemgrenze. Zu beachten wäre bei einem Vergleich mit ТОПОРОВs Ergebnissen, daß wir das Weichheitszeichen bei der Konstituierung der Gruppen aus phonetischen Gründen nicht mitgezählt haben. Daher wird die Verbindung льств (z.B. довольств-

Tabelle 4:

Graphemgruppe	absolute Gesamtfrequenz / in %		Anlaut
TTS	9	0.2	-
TTR	167	3.8	-
TST	361	8.1	-
TSS	3	0.1	-
TSR	104	2.3	-
TSA	8	0.2	-
TRR	5	0.1	-
TAR	3	0.1	-
STT	20	0.5	-
STS	640	14.4	1
STR	1235	27.8	451
SST	112	2.5	17
SSS	2	0.05	2
SSR	576	13.0	1
SSA	5	0.1	-
SRR	11	0.3	1
SAR	2	0.05	-
RTT	11	0.3	-
RTS	27	0.6	-
RTR	548	12.4	1
RTA	64	1.4	-
RST	386	8.7	-
RSS	8	0.2	-
RSR	116	2.6	-
RRT	1	0.02	-
RRR	8	0.2	-
RAT	2	0.05	-
AST	1	0.02	-

во) von ТОПОРОВ zu den fünfgraphemischen gerechnet, während wir sie als viergraphemisch (RSTS) ansehen. Als fünfgraphemische Konsonantenverbindung konnten wir lediglich die Abkürzung РСФСР ermitteln.

Tabelle 5:

Graphemgruppe	absolute Gesamtfrequenz / in %		Anlaut
TSTS	373	46.9	-
TSTR	44	5.5	-
TSST	2	0.25	-
TRTR	1	0.125	-
TRST	3	0.4	-
STST	3	0.4	-
STSR	73	9.2	-
SSTS	7	0.9	-
SSTR	39	4.9	26
RTST	11	1.4	-
RSTS	200	25.2	-
RSTR	35	4.4	-
RSST	1	0.125	-
RRST	1	0.125	-
RAST	2	0.25	-

Anmerkungen

- 1 Vgl. für das Altrussische DIETZE (1975, 39).
- 2 In diesem Zusammenhang soll auch auf die Frequenzangaben zur russischen Fachsprache der Radioelektronik bei ANDREEVA et al. (1965, 50) verwiesen werden.
- 3 Für das Programmieren und die rechentechnische Betreuung schulde ich Frau Chr. Schleiff vom Organisations- und Rechenzentrum der Martin-Luther-Universität in Halle/S. herzlichen Dank.
- 4 Der Klassifikationsmethode von HARARY und PAPER (1957) folgt u.a. auch GVOZDOVIČ (vgl. 1976a, 1976b, 1976c) für das Deutsche.
- 5 Diese Ergebnisse s. bei TOPOROV (1966, 84-116).
- 6 Vgl. dazu TOPOROV (1966, 127-133).

Anhang I: Frequenzen zweigliedriger Konsonantengraphemkombinationen im Russischen

	б	в	г	д	ж	з	и	к	л	м	н	п	р	с	т	ф	х	ц	ш	щ	ni.
б	2																				2548
в	118																				3239
г	1																				908
д	14	209	81	10	10	3	179	879	166	570	83	269	8	1	3	19	7	45	14		2570
ж	105	1		441			2		419												968
з	63	405	15	575			3	44	216	251	511	2	269								2378
и	77						5	1	547	99			338	351	825		67	12	12		2451
к	2	2	41	2	37	383	155	112	54	1463				20	171	78	3	2	136		2661
л	31	14		1			35	15	141	155	33			5							430
м	23	17	535				140	2	2753					31	1033	1537	208	15	40	1	6372
н							3	147	54	109	3826	315	188	4							4650
п	5	122	252	69	190		76	10	2002	506	33	19	71	466	173	41	2	24	60		4121
р	42	532	2	14	11		2251	1209	143	478	1090	254	874	4194	13	86	19	61	124	2	11399
с	31	293		99			420	34	123	673	6	1171	2292	5							5217
т																					267
ф	1	50	1				2	5	1	297			87	55	2						501
х							16														16
ц							79														1632
ш							13														117
щ																					9
n.j	295	1852	411	1790	419	485	3609	5377	3133	10428	1375	7866	4245	7636	1946	327	322	241	478	219	52454

Anhang II: Status zweigliedriger Konsonantengraphemkombinationen im Russischen (nach ALTMANN, G. (1973): Probabilistische Klassifikation von Konsonantenverbindungen des Indonesischen. In: Zeitschrift der Deutschen Morgenländischen Gesellschaft 123, 98-116)

	б	в	г	д	ж	з	к	л	м	н	п	р	с	т	ф	х	ц	ч	ш	щ
б	I	M	I	I	M	P	M	P	M	M	I	P	M	I	I	P	I	I	M	P
в	I	A	I	M	I	P	M	P	A	P	M	M	M	M	I	P	I	I	P	M
г	A	I	I	A	I	I	M	P	M	M	I	P	I	I	I	I	I	P	I	I
д	A	P	P	M	M	M	A	P	A	P	P	M	M	M	M	A	M	P	M	I
ж	P	M	I	P	I	I	M	I	I	P	I	I	I	I	I	I	I	I	I	I
з	P	P	A	P	I	M	M	A	P	P	M	M	I	I	I	I	A	A	I	I
к	I	A	M	I	P	M	M	P	I	M	I	A	P	P	I	I	P	I	M	I
л	M	M	P	M	P	P	M	M	M	P	I	I	M	M	M	I	M	M	P	I
м	P	A	V	M	V	I	A	M	P	P	P	I	M	I	I	V	V	V	I	V
н	I	M	M	P	M	M	M	M	I	P	I	I	M	P	P	I	P	M	M	M
п	I	I	I	I	I	I	M	M	I	M	A	P	M	M	M	I	M	I	I	I
р	M	M	P	M	P	I	M	M	P	M	M	M	M	M	A	P	M	A	P	I
с	M	P	M	M	M	I	P	A	M	M	P	M	A	P	M	P	M	A	P	M
т	A	P	I	M	I	I	P	M	M	M	M	P	P	M	I	I	I	P	I	M
ф	V	I	V	I	V	V	I	M	M	I	I	P	M	M	P	V	V	V	V	V
х	A	P	A	I	I	I	M	M	M	P	I	A	P	M	I	V	V	V	I	V
ц	V	V	V	V	V	V	P	V	V	I	V	V	V	V	V	V	V	V	V	V
ч	I	I	I	I	I	I	M	M	I	P	I	M	I	P	I	I	I	I	P	I
ш	V	A	V	I	V	V	A	P	I	A	A	M	I	M	I	V	V	V	V	V
щ	V	V	V	V	V	V	V	A	P	V	V	V	V	V	V	V	V	V	V	V

Legende: P = "präferiert", d.h., existent und signifikant häufig;
 A = "aktual", d.h., existent und normal häufig;
 M = "marginal", d.h., existent, aber signifikant selten;
 V = "virtuell", d.h., nichtexistent, aber eine "zufällige Lücke";
 I = "inadmissible", d.h., nichtexistent und zugleich eine "strukturelle Lücke".

LITERATUR

ANDREEVA, L.D. et al.

1965 Polučenie pervogo morfologičeskogo tipa russkogo jazyka v pod"jazyke radioelektroniki posredstvom algoritma statistiko-kombinatornogo modelirovanija. In: Statistiko-kombinatornoe modelirovanie jazykov. Moskva-Leningrad, 49-64.

APRESJAN, Ju.D.

1971 Ideen und Methoden der modernen strukturellen Linguistik. Kurzer Abriß. Berlin (Sammlung Akademie-Verlag, Sprache, Bd. 12).

BELONOGOV, G.G., FROLOV, G.D.

1963 Ėmpiričeskie dannye o raspredelenii bukv v rusškoj pis'mennoj reči. In: Problemy kibernetiki 9, 287-305.

DIETZE, J.

1975 Die Sprache der Ersten Novgoroder Chronik. Poznań (Uniwersytet im. Adama Mickiewicza w Poznaniu, Serie Filologia rosyjska, 6).

FRUMKINA, R.A.

1973 Die Anwendung statistischer Methoden in der Sprachforschung. In: Sprachstatistik. Berlin (Sammlung Akademie-Verlag, Sprache, Bd. 22), 272-298.

GWOSDOWITSCH, B.N.

1976a Deutsche Grapheme. In: Wissenschaftliche Zeitschrift der Martin-Luther-Universität Halle-Wittenberg, Gesellschafts- und sprachwiss. Reihe 25/1, 61-81.

1976b Deutsche Graphempaare. In: Wissenschaftliche Zeitschrift der Martin-Luther-Universität Halle-Wittenberg, Gesellschafts- und sprachwiss. Reihe 25/1, 83-90.

GVOZDOVIČ, B.N.

1976c Zakony raspredelenija častot nemeckich grafem. In: Voprosy terminologii i lingvističeskoj statistiki. Voronež, 105 f.

HARARY, F., PAPER, H.

- 1957 Toward a general calculus of phonemic distribution.
In: *Language* 33, 143-169.

KALININA, E.A.

- 1973 Untersuchung lexikostatistischer Gesetzmäßigkeiten auf der Grundlage eines Wahrscheinlichkeitsmodells. In: *Sprachstatistik*. Berlin (Sammlung Akademie-Verlag, Sprache, Bd. 22), 89-143.

KOSCHMIEDER, E.

- 1977 Phonationslehre des Polnischen. München (Münchener Universitäts-Schriften, Reihe der Philosoph. Fakultät, Bd. 13).

OBRATNYJ SLOVAR'

- 1974 *Obratnyj slovar' russkogo jazyka*. Moskva.

ROCZAWSKI, Br.

- 1976 *Zarys fonologii, fonetyki, fonotaktyki i fonostatystyki współczesnego języka polskiego*. Gdańsk.

ŠABROVA, V.A.

- 1978 Sočetačija dvuch soglasnych na morfemnom styke. In: *Russkij jazyk v škole* 4, 85-89.

TOPOROV, V.N.

- 1966 *Materialy dlja distribucii grafem v pis'mennoj forme russkogo jazyka*. In: *Strukturnaja tipologija jazykov*. Moskva, 65-143.

ZUR ÜBERPRÜFUNG DES MENZERATH'SCHEN GESETZES

IM BEREICH DER MORPHOLOGIE

Rainer Gerlach, Göttingen

In der vorliegenden Arbeit soll an deutschem Lexikonmaterial die Hypothese überprüft werden, ob zwischen der Wortlänge (in Phonemen) und der durchschnittlichen Morphemlänge (in Phonemen) der stochastische Zusammenhang:

'Je länger das Wort, um so kürzer seine Morphe.'

besteht; d.h. das Menzerath'sche Gesetz, wie von ALTMANN (1980) abgeleitet, soll im Bereich der Morphologie getestet werden.

Zu diesem Zweck wurde das gesamte Stichwortinventar von WAHRIG (1978) untersucht. Es handelt sich um ca. 16 000 Lemmata. Von jedem Stichwort wurde die Länge, die Zahl seiner Morphe und deren Durchschnittslänge festgestellt. Der phonologischen Transkription der Stichwörter wurde das Phoneminventar des Deutschen von WERNER (1972, 22;40) zugrundegelegt. Die Frikative [s], [x], [ʃ], [ʒ], der velare Nasal [ŋ] sowie [h] und [j] wurden als *ein Phonem* gewertet. Das Problem der mono- bzw. biphonematischen Wertung der Affrikaten und Diphthonge wurde zugunsten des Biphonematischen entschieden (vgl. dazu MORCINIEC 1958 und 1968).

Die Stichwörter in Form von Phonemsequenzen wurden morphologisch analysiert, wobei das Morpheminventar von AUGST (1975) als Basis benutzt wurde. Leider konnte dieses Inventar aber nicht generell verwendet werden, weil es ausschließlich auf Informantenbefragung beruht und daher für linguistische Analysen oft zu ungenau ist. Deshalb wurde ein großer Teil der Stichwörter noch weiter segmentiert, nachdem dies auf Grund eines Substitutionstests angezeigt war (vgl. GREENBERG 1960,

188). Ø-Elemente wurden nicht verwendet; unikale Morphe wurden als eigenständige Segmente behandelt; Eigennamen wurden nicht verwertet; Homonyme bzw. Homographe wurden - je nachdem wie oft diese im Lexikon verzeichnet waren - doppelt oder mehrfach behandelt und gezählt.

Insgesamt wurden 15011 Stichwörter untersucht, die 110672 Phoneme enthielten und sich auf 35346 Morphe verteilten. Die Durchschnittslänge des deutschen Wortes beträgt demnach 7.37 Phoneme (MENZERATH: 7.29), die sich auf durchschnittlich 2.35 Morphe pro Wort verteilen.

Tabelle 1 zeigt die Verteilung der 15011 Einheiten auf die einzelnen Phonem-/Morphtypen; Tabelle 2 verzeichnet dieselbe Verteilung in relativen Zahlen. Hier erkennt man, daß die relativ häufigsten deutschen Wortverbindungen der sechs- und siebenphonemige Zweimorpher und der acht- und neunphonemige Dreimorpher sind (= 40.21 % des Gesamtmaterials). Nimmt man noch die Kombinationen 4/1, 5/1, 5/2, 8/2, 7/3 und 10/3 hinzu, so bilden diese 10 Kombinationen (= 5.95 % der 168 möglichen Kombinationen, bei 21 Phonemen und 7 Morphen) 73.22 % des gesamten Materials.

Die deutsche Gegenwartssprache operiert also in lexikalischer Hinsicht zu ca. 75 % mit nur 10 Phonem-/Morphtypen!

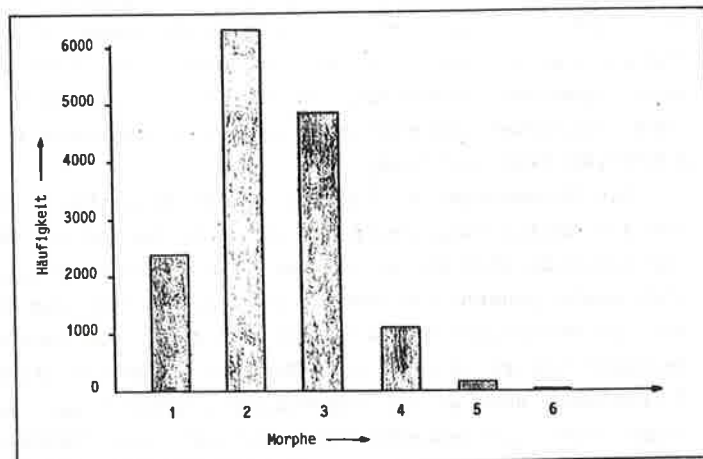


Abbildung 1: Häufigkeitshistogramm der Wortlänge (in Morphzahl)

Abbildung 1 zeigt die Häufigkeitsverteilung auf die einzelnen Morphe; Abbildung 2 veranschaulicht die Häufigkeitsverteilung auf die Phoneme; Abbildung 3 zeigt beide Häufigkeitsverteilungen in integrierter, dreidimensionaler Form.

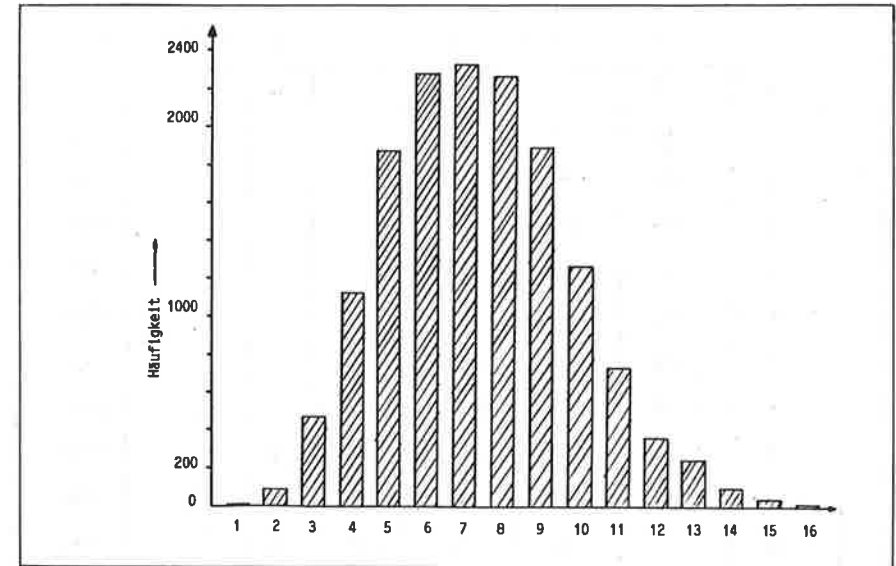


Abbildung 2: Häufigkeitshistogramm der Wortlänge (in Phonemzahl)

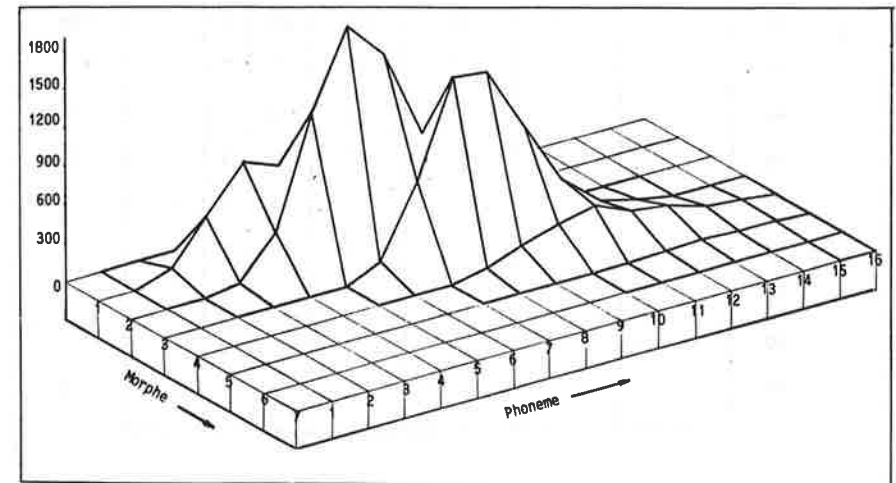


Abbildung 3: Häufigkeitsverteilungen in integrierter (dreidimensionaler) Form

Tabelle 1: Verteilung der 15 O11 Stichwörter auf die einzelnen Phonem-/Morphkombinationen

	Morphe →							Summe
	1	2	3	4	5	6	7	
Phoneme ↑	1	2	3	4	5	6	7	
1	3							3
2	78	2						80
3	422	38						460
4	745	370						1115
5	645	1205	12					1862
6	336	1807	128					2271
7	124	1489	702	7				2322
8	27	812	1380	42	1			2262
9	9	391	1360	142	2			1904
10	2	165	845	242	4			1258
11		55	372	293	16			736
12		8	133	186	14			341
13		1	46	144	26	2		219
14			10	68	20	2		100
15				23	18	3		44
16			1	11	3	3	1	19
17					7	1	1	9
18				1	1	1		3
19						1		1
20								
21							2	2
Summe:	2391	6343	4989	1159	112	13	4	15011

Tabelle 2: Relative Häufigkeitsverteilung der 15011 Stichwörter auf die einzelnen Phonem-Morphkombinationen

	Morphe →							Summe
	1	2	3	4	5	6	7	
Phoneme ↑	1	2	3	4	5	6	7	
1	0.02							0.02
2	0.52	0.01						0.53
3	2.81	0.25						3.06
4	4.96	2.46						7.43
5	4.30	8.03	0.08					12.40
6	2.24	12.04	0.85					15.13
7	0.83	9.92	4.68	0.05				15.47
8	0.18	5.41	9.19	0.28	0.007			15.07
9	0.06	2.60	9.06	0.94	0.01			12.68
10	0.01	1.10	5.63	1.61	0.026			8.38
11		0.37	2.48	1.95	0.11			4.90
12		0.05	0.88	1.24	0.09			2.27
13		0.007	0.31	0.96	0.17	0.013		1.46
14			0.07	0.45	0.13	0.013		0.67
15				0.15	0.12	0.02		0.29
16			0.007	0.07	0.02	0.02	0.007	0.13
17					0.047	0.007	0.007	0.06
18				0.007	0.007	0.007		0.02
19						0.007		0.007
20								
21							0.013	0.013
Summe:	15.93	42.26	33.24	7.72	0.75	0.09	0.027	100%

Um zu überprüfen, ob der vermutete Zusammenhang zwischen Wortlänge und durchschnittlicher Morphemlänge aufweisbar ist, reicht es im Grunde aus, die durchschnittliche Länge jedes Morphtyps in Betracht zu ziehen. Die beobachteten Werte sind in Tabelle 3 angegeben. Die Wortlänge $x = 7$ wurde wegen der kleinen Zahl der Belege ausgeschlossen.

Tabelle 3: Durchschnittliche Morphemlänge in deutschen Wörtern

Wortlänge in Morphen x	1	2	3	4	5	6
Morphdurchschnittslänge y_e	4.53	3.25	2.93	2.78	2.65	2.58

Die durchschnittliche Morphemlänge nimmt also ab, je höhermorphig das Konstrukt wird. Die Abnahmerate ist dabei aber nicht konstant, sondern verringert sich mit zunehmender Morphzahl. Beides veranschaulicht der graphische Abnahmeverlauf (y_e) in Abbildung 4.

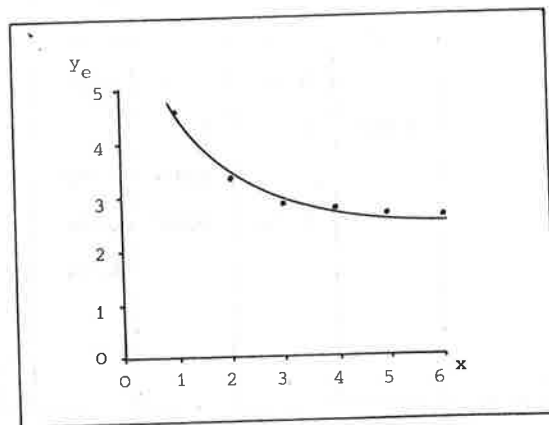


Abbildung 4: Abnahme der Morphemlänge

Der vermutete Zusammenhang zwischen beiden Untersuchungsvariablen ist damit prinzipiell bekräftigt. Es bleibt nur noch zu prüfen, ob dieser Verlauf der von ALTMANN (1980, 3) abgeleiteten Funktion

$$y = A x^b e^{-cx}$$

die das Menzerath'sche Gesetz ausdrückt, entspricht. A , b und c sind Koeffizienten, die vorläufig aus den Daten geschätzt werden müssen (Formeln siehe ALTMANN 1980, 4-5). Aus den obigen Daten ergeben sich folgende Werte für die Konstanten:

$$A = 4.0959$$

$$b = -0.5568$$

$$c = -0.0911$$

Die 'theoretischen' Werte der Morphemlängen (y_t) lassen sich also aus der Funktion

$$y_t = 4.0959 \cdot x^{-0.5568} \cdot e^{0.0911x}$$

berechnen. Der theoretische Funktionsverlauf (y_t) stimmt mit den empirischen Durchschnittswerten (y_e) weitgehend überein. Das zeigt einmal die folgende Übersicht:

Wortlänge in Morphen	1	2	3	4	5	6
y_e	4.53	3.25	2.93	2.78	2.65	2.58
y_t	4.49	3.34	2.92	2.73	2.64	2.61

vor allem aber das Ergebnis des F-Tests (vgl. Anmerkung 1):

$$F_{(2,3)} = 163.55$$

Vergleicht man dieses Ergebnis mit dem tabellierten Wert $F_{0.99(2,3)} = 30.8$, so sieht man, daß die Kurve einen signifikant großen Teil der Datenvariabilität erklärt. Man kann die Gestaltung der deutschen Morpheme als eine Bestätigung des Menzerath'schen Gesetzes betrachten.

Anmerkungen

- 1 Die Koeffizienten müssen aus den Daten geschätzt werden, weil bisher keine theoretischen Ableitungen vorliegen. Der F-Test wird mit der linearisierten Funktion:

$$\ln y = \ln A + b \cdot \ln x - cx$$

durchgeführt. Man nimmt an, daß die Bedingungen für eine multiple Regression erfüllt sind. Da nur mittlere Längen von Interesse sind, verliert man zwar viele Freiheitsgrade und daher auch eine höhere Signifikanz, erspart sich aber auch umfangreiche Rechenarbeit.

ALTMANN, G.

- 1980 Prolegomena to MENZERATH's Law. In: GROTHJAHN, R. (Ed.), Glottometrika 2. Bochum, Brockmeyer, 1-10.

AUGST, G.

- 1975 Lexikon zur Wortbildung, 3 Bde. Tübingen, Narr.

GREENBERG, J.H.

- 1960 A quantitative approach to the morphological typology of languages. IJAL, 26, 178-194.

MORCINIEC, N.

- 1958 Zur phonologischen Wertung der deutschen Affrikaten und Diphthonge. ZPSK, 11, 49-66.
- 1968 Zur Ein- und Zweiphonemigkeit in der deutschen Sprache. Linguistics, 41, 64-74.

WAHRIG, G. (Hg.)

- 1978 dtv-Wörterbuch der deutschen Sprache. München, Deutscher Taschenbuch Verlag.

WERNER, O.

- 1972 Phonemik des Deutschen. Stuttgart, Metzler.

DAS MENZERATHSCHE GESETZ AUF SATZEBENE

R. Köhler, Essen

1. Das von ALTMANN in [1] formulierte und von ihm so getaufte Menzerathsche Gesetz besagt, daß die Länge von linguistischen Einheiten (Komponenten) eine Funktion der Länge des aus ihnen gebildeten sprachlichen Konstrukts ist. Als Differentialgleichung hat das Gesetz die Form:

$$\frac{y'}{y} = -c + \frac{b}{x}, \quad (1)$$

wobei mit x die Länge des Konstrukts und mit y die Länge seiner Komponenten bezeichnet wird. y' ist die erste Ableitung von y .

Die Differenzialgleichung (1) steht für ein *stochastisches* Gesetz; es erfaßt also Tendenzen. Bei der Datengewinnung für empirische Untersuchungen zum Menzerathschen Gesetz müssen daher Mittelwerte verwendet werden. Wir verfahren dabei folgendermaßen:

Eine Stichprobe von N Sprachkonstrukten (Silben, Wörtern, Sätzen ...) besteht aus n_1 Konstrukten der Länge x_1 , n_2 Konstrukten der Länge x_2 , usw. bis n_k Konstrukten der maximalen beobachteten Länge x_k , wobei die Länge als Zahl der unmittelbaren Komponenten angegeben wird (also z.B. Wortlänge als Anzahl der Silben, und nicht als Anzahl der Laute). Die mittlere Länge y_i der Komponenten eines Konstrukts der Länge x_i wird als mittlere Anzahl direkter Bestandteile der Komponenten dieser Konstrukte berechnet:

$$y_i = \sum_{j=1}^{n_i} \frac{y_{ij}}{n_i x_i}; \quad 1 \leq i \leq k, \quad (2)$$

wo y_{ij} die beobachtete Anzahl von direkten Bestandteilen der Komponenten eines x_i -langen Konstrukts ist.

2. Auf der Satzebene folgt aus dem Menzerathschen Gesetz die Hypothese, daß

"in long sentences (measured in number of clauses)
the clauses are shorter and vice versa ..." (i)
(s. [1], S. 9).

Anhand englischer Texte¹ wurde die Hypothese (i) auf folgende Weise empirisch überprüft: Aus den Sätzen der nichtwörtlichen Rede wurde eine zufällige Stichprobe gezogen. Die Satzlänge wurde in Teilsätzen (als Kriterium für einen Teilsatz wurde das Vorhandensein eines finiten Verbs gewählt²), die Teilsatzlänge in Wörtern gemessen. Der Umfang der Stichprobe betrug $N = 843$ Sätze; zwei beobachtete Sätze der Länge 7 wurden wegen zu geringer Frequenz von der Betrachtung ausgeschlossen, so daß $1 \leq x_i \leq 6$ und $N = 841$.

Aus den sechs Einzelverteilungen (vgl. Abb. 1) wurden die Mittelwerte (s. Tab. 1) errechnet. Diese sechs Messungen wurden

Tabelle 1: Mittelwerte des Datenmaterials

x_i	y_i
1	9.7356929779
2	8.3773231506
3	7.3511447906
4	6.7656250000
5	6.1466665268
6	6.2424242424

als empirische Werte für die mittleren Teilsatzlängen in Sätzen mit x_i Teilsätzen verwendet.

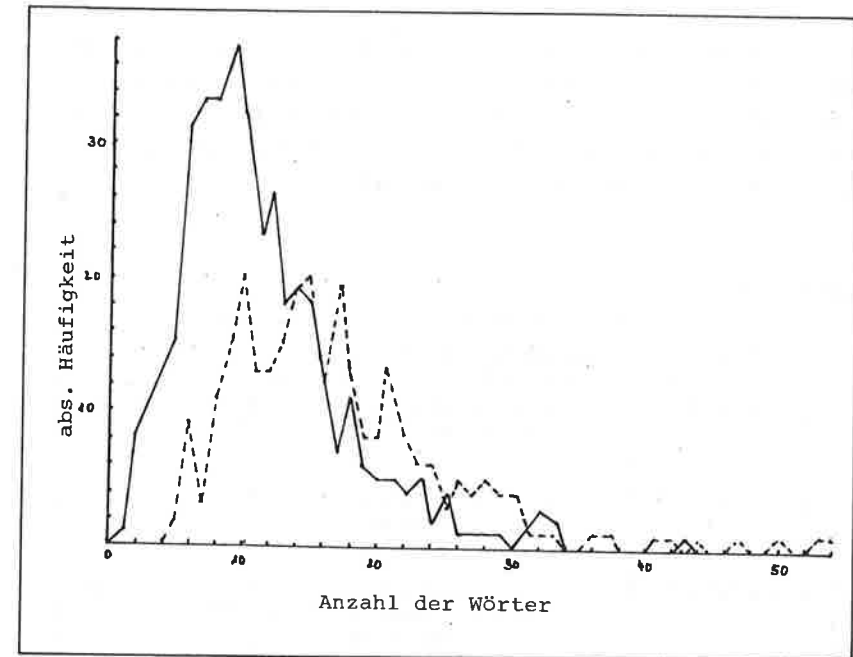


Abbildung 1: Verteilung der Häufigkeit von Teilsatzlängen in Sätzen mit einem (—) und mit zwei (----) Teilsätzen

ALTMANN gibt in [1] die Lösung der Differentialgleichung (1)

$$y = Ax^b e^{-cx} \quad (3)$$

mit den Sonderfällen (4) und (5) für $b = 0$ bzw. $c = 0$ an:

$$y = Ae^{-cx} \quad (4)$$

$$y = Ax^b \quad (5)$$

Obwohl die Vermutung naheliegt, daß auf der Satzebene die monoton fallende hyperbolische³ Funktion (5) zutrifft (Null-Längen kommen nicht vor; auch gibt es keinen Hinweis darauf, daß besonders lange Sätze aus wieder längeren Teilsätzen ge-

bildet werden), wurden alle drei Funktionen nach [1] mit der Methode der kleinsten Fehlerquadrate angepaßt und mit dem F-Test überprüft. Die Ergebnisse (s. Tab. 2) bestätigen die Hypothese (i), aber auch unsere Vermutung, daß (5) gegenüber (3) völlig ausreicht (s. dazu auch Abb. 2).

Tabelle 2:

Funktion	abgeschätzte Parameter	F-Test
$y = e^{\alpha} e^{-cx}$	$\alpha = 2.31595370$ $c = 0.09240026$	$\hat{F}_{1,4} = 48.988$
$y = e^{\alpha} x^b$	$\alpha = 2.28736452$ $b = -0.26885588$	$\hat{F}_{1,4} = 217.81$
$y = e^{\alpha} x^b e^{-cx}$	$\alpha = 2.28939701$ $b = -0.25784883$ $c = 0.00402919$	$\hat{F}_{2,3} = 82.199$
$F_{1,4}(0.05) = 7.709; \quad F_{2,3}(0.05) = 9.552$ $A = e^{\alpha}$		

3. Besteht ein gesetzmäßiger Zusammenhang zwischen zwei oder mehreren Größen, so gilt er natürlich unabhängig davon, in welchen Einheiten diese Größen gemessen werden. Auch linguistische Einheiten wie Satz, Teilsatz, Wort, Morphem usw. sind konventionell festgelegt (bzw. festlegbar). Zur Illustration können Tabelle 3 und Abbildung 3 dienen, die auf den gleichen Daten wie Tabelle 2 und Abbildung 2 beruhen, nur daß neben den finiten Verben jetzt auch Partizipien mit Objekt oder Präpositionalphrase als Teilsatzindikatoren galten. Die Ergebnisse bestätigen wieder die Hypothese (i) und die Funktion (5).

Solange die festgelegten Einheiten konsistent verwendet werden, tangiert ihre Auswahl nicht die Anwendbarkeit von Ge-

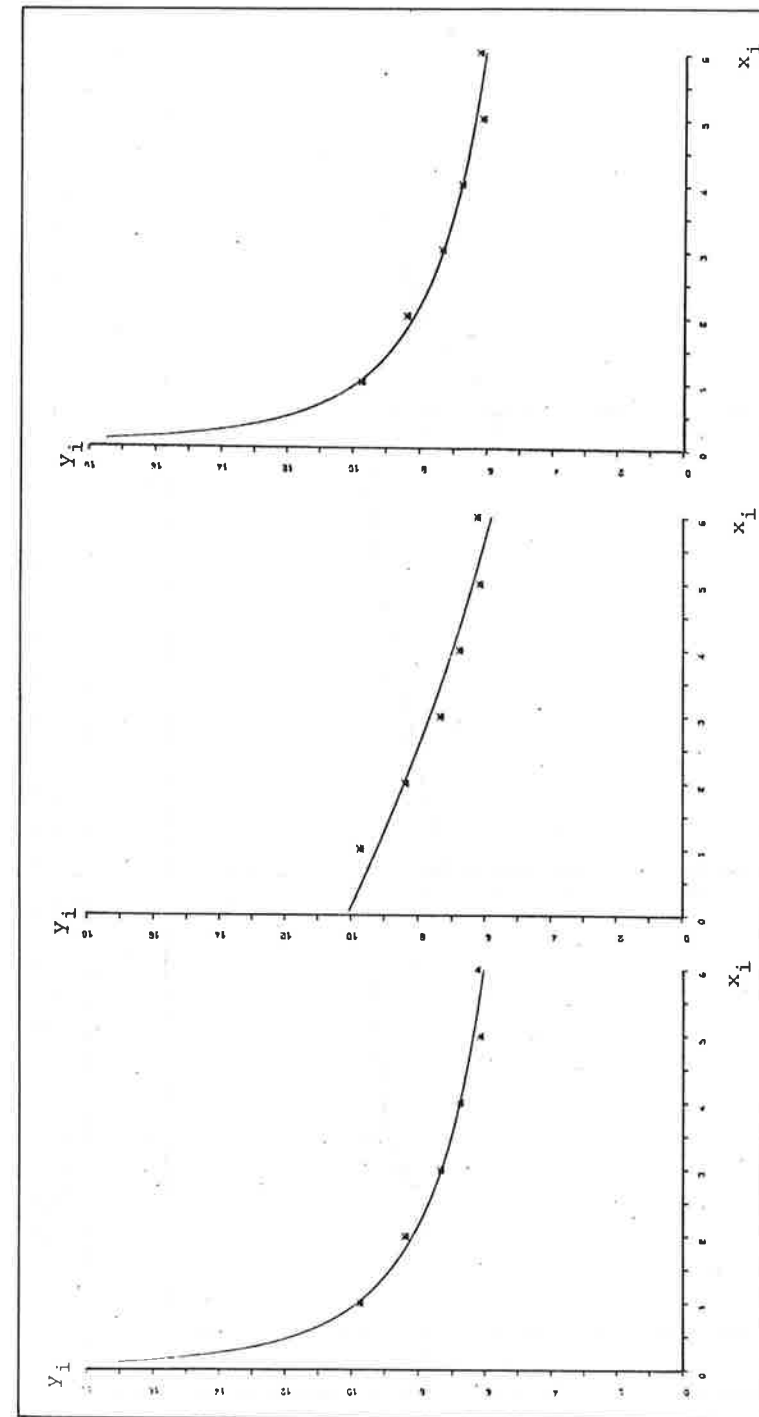


Abb. 2: Stichprobe Crispin; die Graphen der Funktionen (3), (4) und (5) mit den Parametern nach Tab. 2. Die empirischen Werte wurden mit * gekennzeichnet.

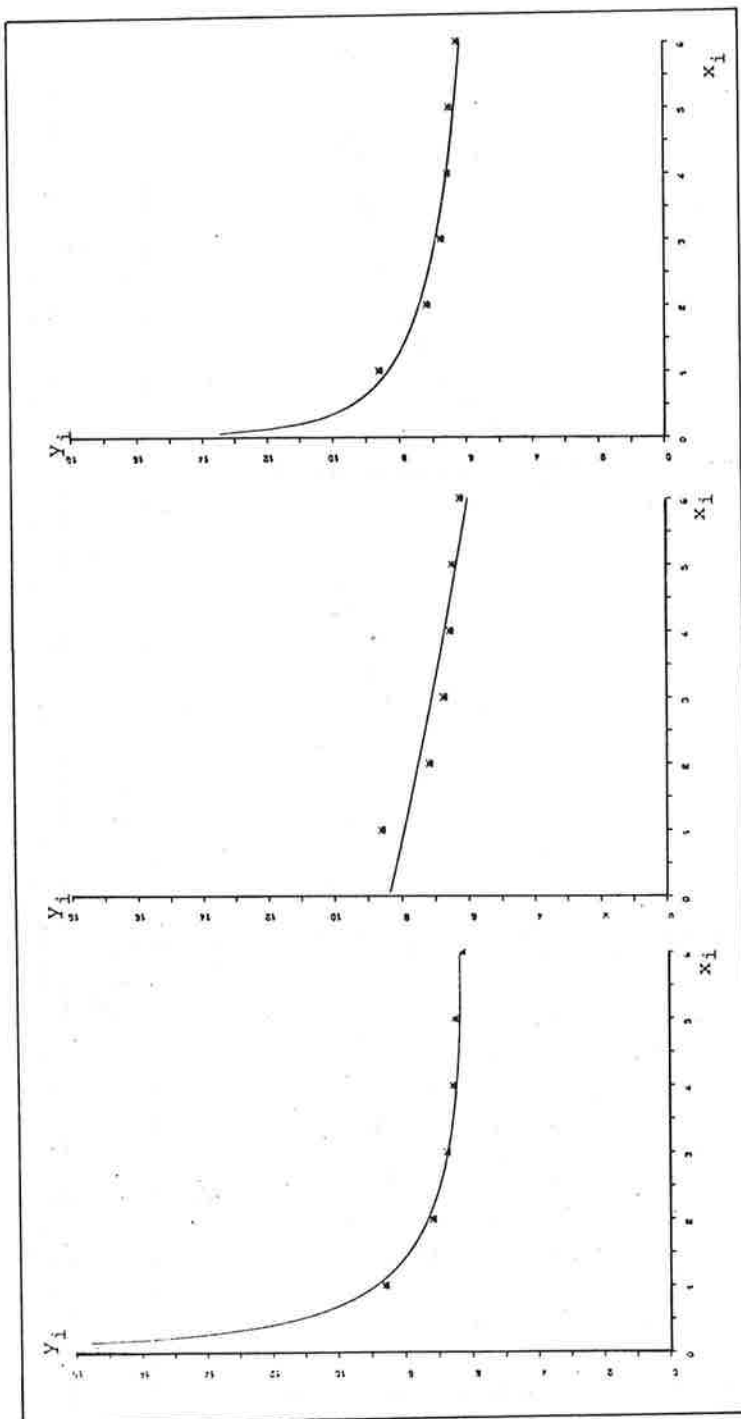


Abb. 3: Stichprobe Crispin; die Graphen der Funktionen (3), (4) und (5) mit den Parametern nach Tab. 3. Die empirischen Werte wurden mit $\times$ gekennzeichnet.

Tabelle 3:

Funktion	abgeschätzte Parameter	F-Test
$y = e^{\alpha} e^{-cx}$	$\alpha = 2.21828914$ $c = 0.05537468$	$\hat{F}_{1,4} = 15.770$
$y = e^{\alpha} x^b$	$\alpha = 2.12151069$ $b = -0.17056614$	$\hat{F}_{1,4} = 75.153$
$y = e^{\alpha} x^b e^{-cx}$	$\alpha = 2.11722649$ $b = -0.37297114$ $c = -0.06659070$	$\hat{F}_{2,3} = 54.177$
$F_{1,4(0.05)} = 7.709; \quad F_{2,3(0.05)} = 9.552$		
$A = e^{\alpha}$		

setzen. Man kann prinzipiell Umrechnungsformeln für die Konvertierung der Einheiten angeben. Es kann aber vermutet werden, daß die Differentialgleichung (1) und die Funktionen (3) bis (5) nicht mehr in dieser Form gelten, wenn in Einheiten gemessen wird, die einer anderen linguistischen Ebene angehören, d.h. wenn Ebenen übersprungen werden (Messung der Satzlänge in Morphemen o.ä.). Ließe sich diese Vermutung empirisch erhärten, könnte dies als Nachweis einer psycholinguistischen Realität der linguistischen Ebenen gelten, da in die Definition der Einheiten ihre Länge in keiner Weise eingegangen ist. Die Befunde von GERŠIĆ und ALTMANN in [3] bezüglich der Abhängigkeit der mittleren Lautdauer von der Wortlänge (in Lautanzahl), wo also die Silbenzahl nicht berücksichtigt wurde, sind bereits eine solche Bestätigung.

4. Eine wesentliche Aufgabe nach der Aufstellung funktionaler Beziehungen wie (3) bis (5) besteht in der *allgemeinen* Bestimmung bzw. in der linguistischen Interpretation der Parameter. Im Fall der Funktion (5), die wir für die Satzebene postulieren, können wir den Parameter α bzw. $A (A = e^{\alpha})$ direkt

interpretieren:

$$A = y_1 \quad (6)$$

Der Parameter A steht für die *primäre Teilsatzlänge*, für die Länge, die der Kürzungstendenz unterliegt; d.h. A ist die mittlere Teilsatzlänge in Sätzen mit nur einem Teilsatz. Mit zunehmender Teilsatzzahl wird von ihr nach unten abgewichen.

Ein Blick auf die Funktionsgleichung von (5) nach Einsetzen der Teilsatzzahl 1 für x_1 zeigt die formale Richtigkeit der Interpretation:

$$y_1 = A \cdot 1^b = A \quad (7)$$

Der Parameter b, der die Steilheit der Kürzung beschreibt, muß weiterhin als empirische Konstante angesehen werden also als sprach- und evtl. textspezifisch.

5. Die bisher gewonnenen Ergebnisse bedürfen weiterer Überprüfung; daher wurden anhand von zwei Texten in deutscher Sprache, ebenfalls auf Satzebene, Daten erhoben⁴. Aus den Anfangskapiteln von U. JOHNSONs "Das Dritte Buch über Achim"⁵ und O. F. BOLLNOWs "Maß und Vermessenheit des Menschen"⁶ wurden die mittleren Teilsatzlängen y_i für die Satzlängen x_i aus den Sätzen errechnet, die keine wörtliche Rede, Verse o.ä. enthalten. Von der Betrachtung ausgeschlossen wurden Sätze der Länge $x_i > 6$ (JOHNSON) bzw. der Länge $x_i > 5$ (BOLLNOW), da sie (pro Satzlänge) weniger als 10 Vorkommen aufweisen ($n_i < 10$). Tabelle 4 und Tabelle 5 enthalten die so gewonnenen Daten, die abgeschätzten Parameter und die Ergebnisse der F-Tests⁷; die Anpassung der Funktion ist in den Abbildungen 4 und 5 dargestellt.

Tabelle 4: (JOHNSON)

x_i	y_i (beobachtet)	$\hat{y}_i$ (berechnet)	n_i
1	9.26443672	9.28773880	329
2	8.57024765	8.42374992	242
3	7.98333263	7.95610332	120
4	7.35087641	7.64013386	57
5	7.35714245	7.40371704	28
6	7.41176414	7.21599007	17
$\alpha = 2.22869539$ $b = -0.14086473$			$N = 793$
$\hat{F}_{1,4} = 68.036 > F_{1,4(0.05)} = 7.709$			

Tabelle 5: (BOLLNOW)

x_i	y_i (beobachtet)	$\hat{y}_i$ (berechnet)	n_i
1	13.95406342	13.50633430	283
2	10.82352829	11.28818893	204
3	10.01626015	10.16360855	123
4	9.16326523	9.43432999	49
5	9.39130402	8.90489197	23
$\alpha = 2.60315895$ $b = -0.25882214$			$N = 682$
$\hat{F}_{1,3} = 48.298 > F_{1,3(0.05)} = 10.128$			

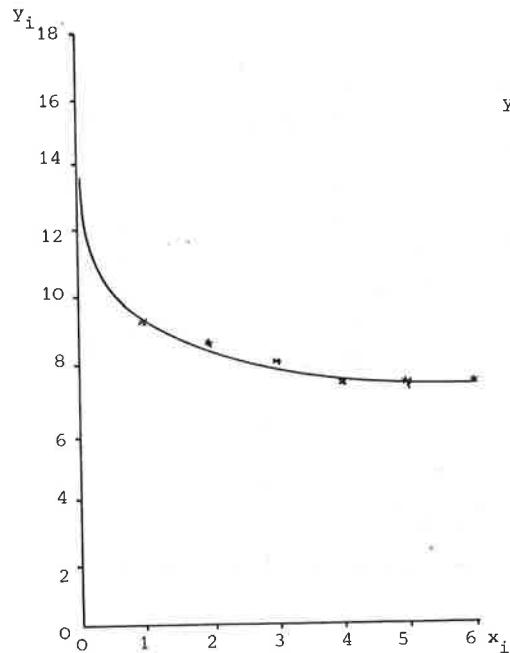


Abbildung 4: Stichprobe JOHNSON;
Graph der Funktion (5) mit
den Parametern aus Tab. 4.

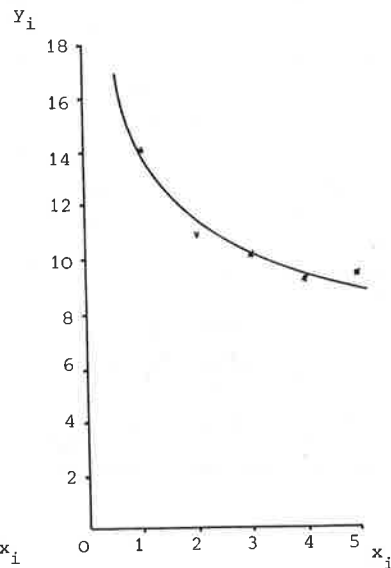


Abbildung 5: Stichprobe BOLLNOW;
Graph der Funktion (5) mit den
Parametern aus Tab. 5.

Die empirischen Werte sind in beiden Abbildungen mit * gekennzeichnet.

Wiederum wurden alle Annahmen bestätigt, womit auch die Interpretation des Parameters A gestützt ist. Der Parameter b erwies sich als textabhängig.

6. Mit der Annahme der Hypothese, daß die Funktion (5) die Kürzungstendenz auf der Satzebene beschreibt, ist die Notwendigkeit verbunden, die Differentialgleichung (1) zu reinterpretieren. ALTMANN ging bei der Aufstellung der Gleichung in [1] davon aus, daß das Glied $-c$ die Kürzung repräsentiert, und das Glied $\frac{b}{x}$ mögliche entgegengesetzte Einflüsse erfaßt. Die Funktion (5) entspricht aber gerade dem zweiten Glied, so daß die vorliegenden Resultate gegen eine konstante Kürzungsrage, we-

nigstens auf der Satzebene, sprechen. Weitere empirische Arbeiten, besonders auf anderen linguistischen Ebenen, scheinen zur Klärung nötig zu sein.

Anmerkungen

- 1 Es wurden die folgenden vier Bände ausgewertet:
CRISPIN (Ed.): Best SF 2 (1956), 3 (o.J.), 4 (1960) und 5 (1971), London.
- 2 Zur Definition des Teilsatzes vgl. [2], S. 27.
- 3 Wegen $b < 0$, da ja Kürzung angenommen wird.
- 4 Für ihre Unterstützung bei der Datenverarbeitung danke ich Frau B. Stöckmann.
- 5 JOHNSON, U., Das dritte Buch über Achim, Frankfurt a.M., 1961.
- 6 BOLLNOW, O.F., Maß und Vermessenheit des Menschen. Philosophische Aufsätze, Göttingen 1962.
- 7 Alle Tests wurden ungewichtet ausgeführt, da die Ergebnisse auch ohne Berücksichtigung der n_i hohe Signifikanz aufwiesen.

LITERATUR

- [1] ALTMANN, G.
1980 Prolegomena to Menzerath's Law. In: GROTHJAHN, R. (Ed.), Glottometrika 2. Bochum, 1-10.
- [2] PIEPER, U.
1979 Über die Aussagekraft statistischer Methoden für die linguistische Stilanalyse. Tübingen.
- [3] GERŠIĆ, S., ALTMANN, G.
1980 Laut - Silbe - Wort und das Menzerathsche Gesetz. In: Forum Phonetikum 21. Frankfurt a.M., 115-123.

PHONOLOGICAL PATTERNING IN ENGLISH AND GERMAN VERSE: A COMPUTER-ASSISTED APPROACH

David H. Chisholm, Arizona

Introduction

A few years ago I investigated the sound repetition patterns in a corpus of nineteenth century German iambic pentameter verse with the aid of two computer programs written in SNOBOL 4. What distinguished this investigation from previous phonological studies of German verse was a procedure which takes into account not only conventional manifestations of phonological recurrence, such as end-rhyme, alliteration and assonance, but *all* recurrence patterns in a given poem or prose selection. This procedure treats all patterns of equivalence in binary terms, and distinguishes recurrence pairs according to the size of the gap between the two syllables in which the sound occurs. The analysis of the German corpus, which has since been expanded to include early twentieth century as well as nineteenth century poets, yielded some striking results which distinguish the sound patterning of metered verse, free verse and prose, phonostylistic characteristics of the work of individual poets and diachronic changes in nineteenth and twentieth century verse.

The fruitfulness of this approach to German verse suggested its application to English verse of roughly the same period to determine to what extent the predominant phonological patterns of English verse differ from, or are similar to, those which prevail in German. Application of this type of computer-assisted approach to English verse also enables the investigator to isolate those phonological patterns which distinguish the verse of individual poets and periods. Moreover, it provides a reliable basis for evaluation of literary studies which discuss the sound structure of English verse.

1. Phonological Frames

The procedure, which has been described in CHISHOLM (1976), may be summarized as follows: Pairs of equivalent sounds are distinguished on the basis of proximity by assigning each pair to the appropriate *phonological frame*. Equivalence between contiguous syllables is assigned to Frame One, between alternate syllables to Frame Two, and so forth. In each phonological frame, all pairs of syllables are scanned for equivalence. In the line given in (1), the three instances of prevocalic /m/ generate *three pairs* of equivalent sounds which are assigned to different frames:

(1) Maiden and mistress of the months and stars

Frame 3: Mai-/ mis-
Frame 4: mis-/ months
Frame 7: Mai-/ months

In this way, patterns of equivalence are identified in binary terms within a positional framework. One restriction should be noted: equivalence is recognized only between *corresponding parts* of syllables; in the sequence "maiden and," for example, the postvocalic /n/'s constitute an equivalence, while the two /d/'s, one being prevocalic and the other postvocalic, are not recognized as a pair. Each poem is viewed as a single long string of syllables, so that the two members of a pair may be located in different lines. The patterning in the German poems attests to the theoretical soundness of distinguishing *vertical equivalence between lines* from all other equivalence. The frames which include vertical equivalence are called *primary frames*; all other frames are *secondary frames*. Primary Frame 1 includes all vertical correspondence between successive lines, Primary Frame 2 all vertical correspondence between alternate lines, and so forth. This distinction, of course, is meaningful only for verse, in which the "line" functions as a rhythmic entity.

2. The Data

The corpus of English iambic pentameter verse selected for this investigation consists of forty sonnets by four nineteenth century English poets - Wordsworth, Keats, Rossetti and Swinburne - and 148 lines of blank verse from Swinburne's *Atalanta in Calydon* which is divided into ten fourteen- and sixteen-line segments. The data thus provide a fairly representative sampling of English iambic pentameter verse from both the Romantic period (Wordsworth, Keats) and the Victorian period (Rossetti, Swinburne). The inclusion of the blank verse passages from *Atalanta in Calydon* (lines 1-64 and 121-204) enables us to compare rhymed and unrhymed iambic pentameter verse. Furthermore, it is the text which B.F. SKINNER (1941) used in support of his theory of "formal perseveration" and it can therefore be used as a means of testing that theory. A corpus of fifty nineteenth century German sonnets by Goethe, Schlegel, Brentano, Heine and Mörike is included for a general comparison of phonological patterning in nineteenth century German and English verse.

3. Transcription of the Data

The phonemic transcription of the English data follows the phonetic description of English Received Pronunciation given by Daniel JONES (1956) and A.C. GIMSON (1962). The forty-five phonemes of English described by JONES and GIMSON are given in (2) with examples of words in which they occur, and the keypunch character used for each phoneme. In this investigation the diphthongs /ei, ai, oi, ou, au/ and the affricates /tʃ, dʒ/ are treated as single phonemes. The transcription of words such as *tire*, and *tower*, which contain a diphthong followed by schwa, depends on their syllabic relation to the underlying meter: if the sequence /aie/ or /aue/ occupies two metrical positions it is assumed to be disyllabic and is transcribed accordingly; if it occupies only a single metrical position, it is treated as monosyllabic and reduced to the diphthongal glide /a:ə/, with

(2) Phonemes of British English

a)

Vowel Phonemes	Example	Keypunch Character
/i:/	see	1
/ɪ/	bit	I
/e/	get	E
/æ/	cat	A
/ɑ:/	father	6
/ɔ/	hot	O
/ɔ:/	saw	8
/u/	put	U
/u:/	too	7
/ʌ/	but	^ (upper case Q)
/ə:/	bird	9
/ə/	china	ə
/ei/	day	EI
/ai/	die	AI
/oi/	boy	OI
/ou/	go	OU
/au/	how	AU
/ie/	here	Iə
/eə/	there	Eə
/uə/	tour	Uə
/a:ə/	tire, hour	6ə (monosyllabic form of /aie, aue/)

(N.B. ə corresponds to α in (3) b) and to ö in Figure 1(cont.).)

b)

Consonant Phonemes	Example	Keypunch Character
/p/	pole	P
/b/	bowl	B
/t/	toll	T
/d/	dole	D
/k/	coal	K
/g/	goal	G
/tʃ/	cheap	C
/dʒ/	jeep	X
/f/	feel	F
/v/	veal	V
/θ/	thigh	Q
/ð/	thy	J
/s/	seal	S
/z/	zeal	Z
/ʃ/	vicious	(
/ʒ/	vision	)
/h/	heal	H
/m/	meat	M
/n/	neat	N
/ŋ/	sing	/
/l/	leap	L
/r/	reap	R
/w/	weep	W
/j/	year	J

the result that *tire* and *tower* become homophones. This approach is consistent with GIMSON's interpretation of /aiə/ and /auə/ and the levelling of /a:ə/ and /ɑ:ə/ into a single phoneme (1962: 133).

Stress is treated as a binary opposition + STRESS/-STRESS.

+ Stress is assigned to:

- (1) all syllables bearing primary stress in polysyllabic words and all monosyllabic nouns, non-auxiliary verbs and adjective
- (2) all syllables bearing primary stress *relative to their immediate morphological context* in words of three or more syllables (e.g. *grandfather*, *apple tree*, *workmanship*, *overhear*).

In the input data, all stressed syllables are preceded by a plus sign.

The *syllable* may be identified as a vowel nucleus optionally preceded and/or followed by one or more consonants; when two vowels are separated by a cluster of consonants, however, it is often difficult to determine where one syllable ends and the next begins. GIMSON gives as an example the word *extra* /'ekstrə/, in which the phoneme /s/ could belong either to the first or second syllable and still be in accord with the phonological rules of the English language. For computer-assisted analysis syllable-boundary is determined in the following necessarily somewhat arbitrary way. It is located:

- (1) between words (e.g. *first day*)
- (2) immediately before all free lexical items (e.g. *first-born*) and derivational suffixes beginning with a consonant (e.g. *first-ly*)
- (3) in the middle of consonant clusters where neither (1) nor (2) apply. Such clusters are divided evenly, with any extra consonants placed to the right of the syllable boundary (e.g. /'eks-trə, 'læs-tɪŋ, 'raɪ-dɪŋ, ɪk-'spænd/).

If word-final orthographic *r* is followed by a word-initial vowel (so called "linking-r"), it is transcribed as a post-vocalic /r/ (e.g. *far away* /fɑ:r ə'weɪ/ and the word-boundary serves as the syllable boundary.¹ In the input data, syllable-boundaries are marked by an asterisk. As an illustrative example, the machine-readable transcription of a sonnet by Keats is given in (3).

(3) John Keats, "Keen, fitful gusts ..."

a) Orthographic Representation:

Keen, fitful gusts are whispering here and there
 Among the bushes half leafless, and dry;
 The stars look very cold about the sky,
 And I have many miles on foot to fare.
 Yet feel I little of the cool bleak air,
 Or of the dead leaves rustling drearily,
 Or of those silver lamps that burn on high,

Or of the distance from home's pleasant lair:
 For I am brimful of the friendliness
 That in a little cottage I have found;
 Of fair-haired Milton's eloquent distress,
 And all his love for gentle Lycid drowned;
 Of lovely Laura in her light green dress,
 And faithful Petrarch gloriously crowned.

b) Phonological Transcription:

```
&JOHN KEATS 06, PAGE 297B "KEEN, FITFUL GUSTS..."
%
1. +KIN, +FIT*FUL +G*STS 6 +WIS*PRI/ HI* AND JE*
2. **M*/ J* +BU*(IZ, +H6F +LIF*LI* AND +DRAI;
3. J* +ST6Z +LUK +VE*RI +KOULD **BAUT J* +SKAI,
4. AND AI HAV +ME*NI +MAILZ ON +FUT T7 +FE*,
5. JET +F1L AI +LI*T*L OV J* +K7L +BLIK +E*,
6. 8 OV J* +DED +LIVZ +RUS*LI/ +DRI*RI*LI,
7. 8 OV IOUZ +SIL*V* +LAMPS IAT +B9N ON +HAI,
8. 8 OV J* +DIS*T*NS FROM +HOUMZ +PLE*S*NT +LE*,
9. F8 AI AM +BRIM*FUL OV J* +FRIEND*LI*NIS
10. IAT IN * +LI*T*L +KO*TI* AI HAV +FAUND;
11. OV +FE*HE*D +MIL*T*NZ +E*LD*K*NT DIS*+TRES
12. AND 8L HIZ +L*V F8 +XEN*T*L +LI*SID +DRAUND,
13. OV +L*V*LI +L8*R* IN H9 +LAIT +GRIN +DRES,
14. AND +FEIQ*FUL +PI*TR6K +GL8*RI*+S*LI +KRAUND.
```

The transcription of the German texts, based on Moulton's *Sounds of English and German*, Siebs' *Deutsche Aussprache*, and the Duden *Aussprachewörterbuch*, has been described elsewhere (CHISHOLM 1976: 7).

4. The Programs

The two SNOBOL 4 programs initially used for German verse and prose may be utilized for English by simply changing the inventory of consonant and vowel phonemes. (In this way the programs can be used for phonological studies of virtually any language.) The *Phoneme Frequency Program* computes the absolute and relative frequency of all sounds in prevocalic, vocalic,

and postvocalic position in stressed and unstressed syllables. The *Phonological Frame Program* identifies all recurrence-pairs in a given number of frames and prints the results in graphic and statistical form. A sample of the output of the *Phonological Frame Program* is given in *Figure 1* (a and b), which shows the graphic representation and statistics for Frame Two of the transcribed Keats sonnet in (3). The graphic representation shows only those syllables which contain equivalences in this frame. Comparison of the graphic representations of the different frames of a poem gives a general impression of its phonological development from beginning to end. The statistics, which can be derived from the graphic representation, show the frequency of sound repetition pairs in stressed and unstressed syllables (labelled S and U) in prevocalic, vocalic and postvocalic position. The amount of phonological equivalence in each frame is expressed as a ratio of the number of sound repetition pairs to the numbers of syllable pairs in that frame.

5. The Consonant-Vowel Ratio in German and English

Table 1, based in part on the output of the frequency program, gives for both English and German the absolute and relative frequency of vowels and consonants in (a) the phoneme inventory and (b) in stressed and unstressed syllables in nineteenth century verse. Since the phonemic inventory of English is substantially larger than that of German, we can expect *less phonological recurrence* in English. Note also that English, due largely to the loss of the unstressed vowel /ə/ in the thirteenth and fourteenth centuries and concomitant increase in lexical monosyllables, has a much lower proportion of *unstressed syllables* than German. Furthermore, due in part to the same type of historical process (loss of unstressed inflectional morphemes) English has a lower proportion of unstressed *consonants* than German. Comparison of the English sentence "We stay with our old friends" with its German equivalent "Wir bleiben bei unseren alten Freunden" (unstressed syllables underlined)

Table 1: Vowels and Consonants in English and German Verse

A. Phonemic Inventory of the Language

	English		German	
Vowels	21	46.67 %	18	47.37 %
Consonants	24	53.33 %	20	52.63 %

B. Frequency in Verse

	Four English Poets		Five German Poets	
Stressed Vowels	2380	32.66 %	2774	34.52 %
Stressed Consonants	4908	67.34 %	5262	65.48 %
Stressed Phonemes	7288	100.00 %	8036	100.00 %
Unstressed Vowels	3254	40.62 %	4807	36.53 %
Unstr. Consonants	4757	59.38 %	8351	63.47 %
Unstressed Phonemes	8011	100.00 %	13158	100.00 %
Total Vowels	5634	36.83 %	7581	35.77 %
Total Consonants	9665	63.17 %	13613	64.23 %
Total Phonemes	15299	100.00 %	21194	100.00 %
Total Stressed Syllables	2380	42.24 %	2774	36.59 %
Total Unstressed Syllables	3245	57.76 %	4807	63.41 %
Total Syllables	5634	100.00 %	7581	100.00 %

illustrates the difference. The English sentence has three unstressed syllables containing three consonant phonemes, while the German equivalent has seven unstressed syllables containing *twelve* consonant phonemes. In stressed syllables, on the other hand, English shows a higher proportion of consonants than does German. In the sentence given above, the three stressed English syllables contain nine consonant phonemes, while the four stressed German syllables contain only seven. This is reflected in verse, where English has a higher proportion of stressed consonants than German.

6. Phonological Patterning in English Verse and Prose

Table 2 (and the accompanying graph in Figure 2) gives the frequency and distribution of phonological equivalence in Frames 1-10 of forty sonnets by Wordsworth, Keats, Rossetti and Swinburne, ten 14-16 line segments from the blank verse passages of Swinburne's tragedy *Atalanta in Calydon*, and a 420-syllable prose sample from his essay on William Blake. The frequency and distribution of fifty nineteenth century German sonnets and twenty-five fourteen-line segments from Goethe's blank verse drama *Iphigenie auf Tauris* are included for further comparison. In the four verse corpora the even frames (2, 4, 6, 8, 10) are consistently higher than the odd frames (1, 3, 5, 7, 9) which means that there is a significantly greater amount of equivalence between two metrically strong or two metrically weak positions in the iambic pentameter line than between a strong and a weak position. This contrast between even and odd frame equivalence reflects the principle of alternation which inheres in both English and German iambic verse. Frame Ten (which includes rhymes in successive lines) is higher than the other even frames in the German and English sonnets. In the *unrhymed* blank verse samples, Frame Ten is significantly higher in German (with Frame Six in second place); in English, on the other hand, Frames 6, 4 and 2 are higher than Frame Ten. This reflects a major difference in the extent of vertical patter-

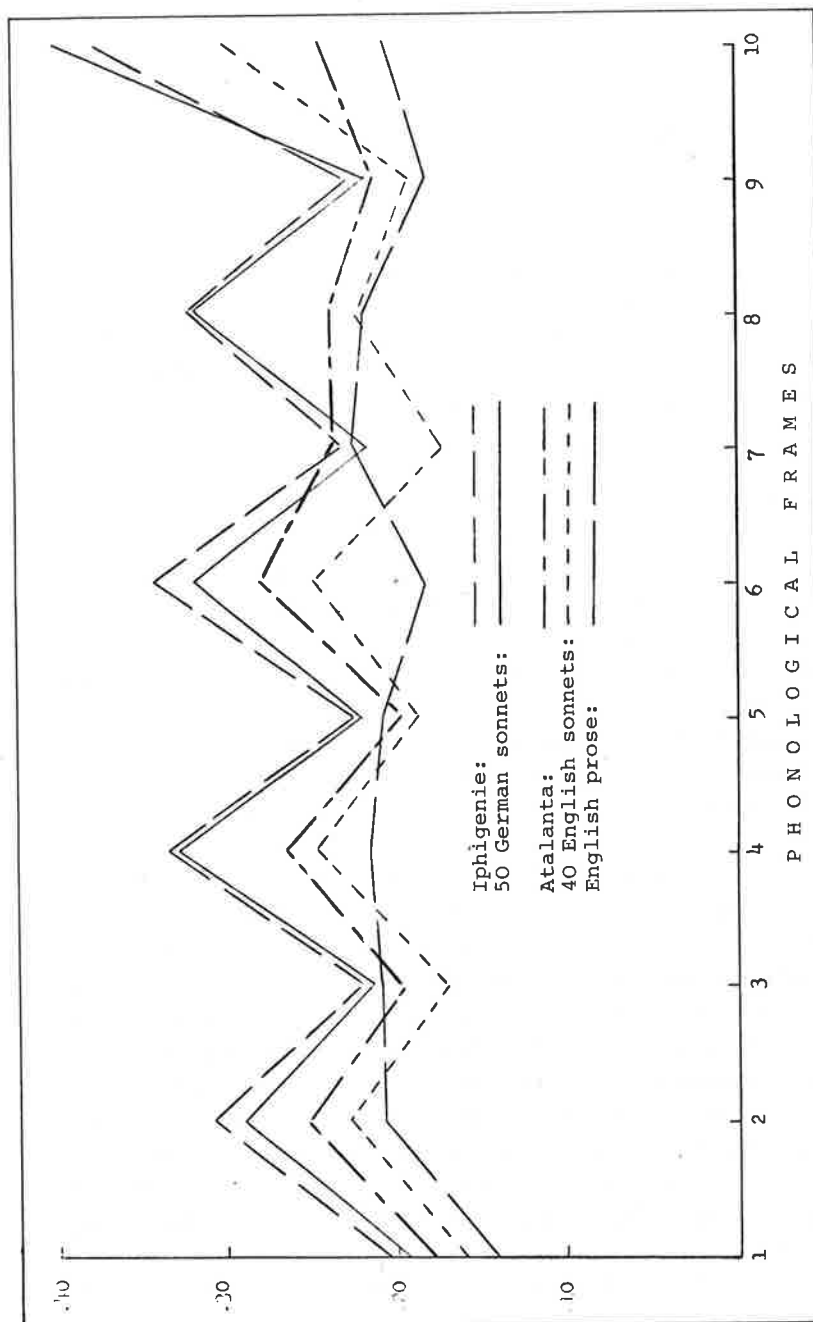


Figure 2: Ratio of Sound Pairs to Syllable Pairs in Frames 1-10 in English Poetry and Prose

Table 2: Ratio of Sound Pairs to Syllable Pairs in Frames 1-10 in German and English Iambic Pentameter Verse and a Sample of English Prose

FRAMES

	1	2	3	4	5	6	7	8	9	10	Total
GERMAN											
Five Poets	.198	.292	.217	.330	.220	.320	.219	.320	.221	.398	.273
<i>Iphigenie</i>	.205	.310	.221	.334	.223	.343	.235	.321	.226	.375	.279
ENGLISH											
Four Poets	.165	.229	.173	.247	.189	.244	.174	.225	.193	.299	.213
<i>Atalanta</i>	.182	.257	.199	.265	.197	.282	.237	.239	.213	.245	.232
Prose	.141	.208	.209	.218	.210	.182	.226	.260	.183	.208	.204

ning in English and German verse: in German verse, vertical patterning predominates *throughout the verse line*, so that unrhymed as well as rhymed iambic pentameter verse consistently shows more equivalence in Frame Ten (= Primary Frame One) than in Frames 2, 4, 6 and 8. The English blank verse of Swinburne's *Atalanta* on the other hand, does not show this tendency.

A measure of equivalence for comparing one poem or group of poems with another may be obtained by dividing the total number of sound repetition pairs in a given number of frames by the total number of syllable pairs in those frames. This measure is called the *coefficient of phonological equivalence*. Computation of the total equivalence in Frames 1-10 in the blank verse and the sonnets shows that this coefficient is somewhat *greater* in the blank verse passages than in the sonnets for both English (.299 versus .213) and German (.279 versus .273)². In contrast to the graphs of the iambic pentameter poems, the graph of the prose sample shows no regular pattern

of alternation. Table 2 and Figure 2 do show one feature that both the verse and prose samples share in both languages: there is a marked tendency - particularly in prose - to avoid equivalence in Frame One. This seems to reflect a feeling among writers of English and German that excessive phonological equivalence between contiguous syllables is "bad style" and should be avoided³.

7. Phonological Frames and Skinner's "Formal Perseveration"

In an article entitled "A Quantitative Estimate of Certain Types of Sound Patterning in Poetry", the behaviorist B.F. SKINNER argues that sound repetition in verse is merely an example of what he calls "formal perseveration" or "formal strengthening". He describes this principle by saying that "when a speech-sound is once emitted, the probability of emission of that sound is temporarily raised" (1941: 64), but then falls to the chance level after a certain number of intervening syllables. To support this general claim, SKINNER analyzed alliteration and assonance in the first 500 iambic pentameter lines of Swinburne's "Atalanta in Calydon". He restricted his investigation to accented syllables, thereby omitting over 40 % of the linguistic material, and claimed that "the omitted material may be regarded either as not contributing in any effective way to the sound pattern or as not permitting any variation in the behavior of the writer" (p. 65). It is easy to demonstrate, however, that unaccented syllables do contribute to the sound pattern and that they also allow - albeit to a lesser degree than accented syllables - stylistic variation on the part of the poet. Thus SKINNER prejudiced his results from the outset by excluding from consideration over 2000 of the 5000 syllables in his sample. Furthermore he limited his investigation of alliteration to only ten of the 24 consonant phonemes of English. The first part of his investigation merely shows that here is a lot of line-internal alliteration in Swinburne's blank verse, which is hardly a startling fact. The

second part, despite unreliability of results, is methodologically of considerable interest. In order to compute the number of syllables between alliterating sounds, SKINNER placed the initial consonants of all selected syllables in order on a long tape:

"The tape was then passed under cards in which windows were cut to display consonants with any desired number of intervening syllables, beginning with adjacent pairs. The numbers of pairs of each consonant at distances up to that of eight intervening syllables were tallied by moving the type under each card one space at a time throughout its length." (p. 67 f.).

This procedure is essentially the same as that of the previously mentioned Phonological Frame Program, the difference being that the latter takes *all* syllables (both stressed and unstressed) and *all* phonemes (prevocalic, vocalic and postvocalic) into consideration. The essential weakness of SKINNER's procedure is that the size of the gaps between repeated sounds in his data bears no constant relationship to that of the original poem, as illustrated by three lines from *Atalanta in Calydon* given in (5):

	1	2	3	4	5	6	7	8	9
	For	if	<i>sleep</i>	have	no	<i>mercy</i>	and	<i>mans</i>	<i>dreams</i>
(5)	11	12	13	14	15	16	17	18	19
	<i>Bite</i>	to	the	<i>blood</i>	and	<i>burn</i>	into	the	<i>bone</i>
	21	22	23	24	25	26	27	28	29
	What	<i>shall</i>	this	<i>man</i>	do	waking?	By	the	<i>gods</i>

Since SKINNER's long tape contains only the underlined syllables, his pairs having "0 intervening syllables" would have 0, 1, 2 and 3 intervening syllables in the original poem (Frames 1, 2, 3, 4), those having "1 intervening syllable" would have 1, 3, 4 and 5 (Frames 2, 4, 5 and 6), those having "2 intervening syllables" would have 4, 5, 6, 7 and 8 (Frames 5, 6, 7, 8, 9) and those having "8 intervening syllables (the largest gap considered by SKINNER) would have 19 and 20 intervening syllables (Frames 20, 21). The Phonological Frame Program gi-

ves a much more accurate picture of the distribution of sound patterns in verse and prose, since no part of the text is excluded from consideration.

Even if SKINNER had taken all syllables into account, however, there is still a fundamental flaw in his reasoning. In *Atalanta*, Swinburne is consciously making extensive use of the *conventional device* of alliteration. In order for the listener to perceive the alliterative pattern, the repeated sound must occur before the impression created by its antecedent is blotted out by the intervening phonological material. In isolation the conventions of alliteration and assonance, which require repetition of only a single phoneme, appear to be consciously perceivable as such over a span of not much more than about ten syllables, roughly the length of the iambic pentameter line. The convention of end-rhyme, on the other hand, which is positionally bound and requires the repetition of more phonetic material (usually two or more phonemes), can be perceived over a span of 2 or 3 *lines* (as in the sonnet form). Similarly, the repetition of the *refrain* in stanzaic forms may be separated from its initial occurrence by many lines, since the repeated event consists of a span of phonemes (and words) repeated in the same sequence.

All of these and similar conventions, however, are not examples of "formal perseveration" as defined by SKINNER. To determine whether sound repetition patterns in poetry reflect what SKINNER calls "formal perseveration", we need to consider non-conventional as well as conventional patterns of phonological equivalence. The Phonological Frame Program does this by taking all phonemes into consideration.

One way of determining whether "formal perseveration" is an operative principle in English verse is to compare the amount of phonological equivalence in each frame. If SKINNER's assertion is correct, we should expect the greatest amount of equivalence in Frame Two, and successively less in Frames 4, 6, 8, 10 etc.; similarly, we would expect a gradual decrease in equivalence in odd frames after Frame One. This assumption was tested by computing the equivalence in each of the first

Table 3: Highest Coefficients of Equivalence in 148 Iambic Pentameter Lines of *Atalanta in Calydon*

A. Highest Coefficients of Frames 1-20

Coefficient of Equivalence	Frame	Lines of Text
.399	2	17- 22
.373	10	49- 63
.365	4	49- 64
.350	20	1- 16
.336	20	17- 32
.328	6	163-176
.325	6	49- 64
.321	4	17- 32
.318	6	1- 16
.316	14	1- 16
.312	20	163-176
.311	9	1- 16
.311	12	1- 16
.309	8	33- 48
.308	7	135-148
.304	2	33- 48
.300	10	135-148

B. Highest Coefficients of Equivalence in Odd Frames 1-19

Coefficient of Equivalence	Frame	Lines of Text
.31	9	1- 16
.31	7	135-148
.29	13	17- 32
.29	11	49- 64
.27	13	163-176
.26	15	33- 48
.25	13	149-162
.25	1	191-204
.24	7	121-134

twenty frames in the first 148 unrhymed iambic pentameter lines of Swinburne's *Atalanta in Calydon*. As mentioned previously, the text was divided into ten segments to facilitate comparison with the sonnets. Table 3 lists in rank order those frames which show the highest equivalence (from .399 - .300) in these segments. (The frames with less equivalence than those included here range from .101 - .292.) This data offers strong counter-evidence to SKINNER's theory of "formal perversion", since the frames which occur most frequently in this group are Frames 20 and 6, followed by 4, 2 and 10. If SKINNER's theory were correct, there would be numerous occurrences of Frames 2 and 4 at the top of the list, and Frames 10 and 20 would not even occur in this range. Taken separately, the *odd* frames yield the same kind of evidence; Frame 13 is the highest of the odd frames in three of the segments (equivalence ratio = .29, .27, .25) and Frame 7 in three segments (.31, .24, .19) while Frames 9, 11, 15 and 1 are each highest in only one segment and Frames 3 and 5 are not highest in any segment. Table 2 as well indicates that SKINNER's general assumption is incorrect: in the first ten frames of the forty English sonnets, Frames 10, 4 and 6 each have more equivalence than Frame 2. In the fifty German sonnets and in the twenty-five segments from *Iphigenie*, Frame 2 shows the lowest equivalence of the first five even frames. Nor does SKINNER's assumption hold for Swinburne's *prose*; among the first twenty frames of the prose sample from his essay on William Blake, Frames 8, 7, 15, 4 and 13 show the greatest amount of equivalence.

8. Coefficients of Equivalence in Even and Odd Frames

The graph in Figure 2 emphasizes the importance of the distinction between even and odd-numbered frames; for phonological equivalence in even frames tends to support the underlying meter, while odd-frame equivalence tends to disrupt it. Table 4 gives the coefficients of equivalence in even and odd frames for the first ten frames in the forty sonnets by Wordsworth,

Table 4: Coefficients of Equivalence in the Sonnets of Wordsworth, Keats, Rossetti and Swinburne

A. Even Frames (2, 4, 6, 8, 10)

	Sound Pairs	Syllable Pairs	Coefficient of Equivalence
Swinburne	1855	6700	.277
Rossetti	1649	6700	.246
Keats	1631	6700	.243
Wordsworth	<u>1524</u>	<u>6700</u>	<u>.227</u>
Total	6661	26800	.249

B. Odd Frames (1, 3, 5, 7, 9)

Rossetti	1306	6750	.193
Swinburne	1242	6750	.184
Wordsworth	1191	6750	.176
Keats	<u>1096</u>	<u>6750</u>	<u>.162</u>
Total	4835	27000	.179

C. Frames 1-10

Swinburne	3097	13450	.230
Rossetti	2955	13450	.220
Keats	2727	13450	.203
Wordsworth	<u>2715</u>	<u>13450</u>	<u>.202</u>
Total	11494	53800	.214

D. Difference between Even and Odd Frames

Swinburne	.093
Keats	.081
Rossetti	.053
Wordsworth	<u>.052</u>
Total	.070

Keats, Rossetti and Swinburne, as well as the difference between even- and odd-frame equivalence for each poet. Swinburne is distinguished from the other English poets by a very high coefficient of equivalence in even frames, and Keats is distinguished by a low coefficient in odd frames. The *total* equivalence in Frames 1-10 clearly distinguishes the Victorian poets Swinburne and Rossetti from the Romantic poets Wordsworth and Keats. The contrast between even and odd frames is greatest in Swinburne and Keats and lowest in Rossetti and Wordsworth. A poem which shows an extremely high degree of contrast between even and odd-frame equivalence is Keats' sonnet "Great spirits now on earth ...". Figure 3 shows the distribution of equivalence in the first ten frames of this poem compared with that in all forty English sonnets. On the other end of the scale are those poems in which the contrast between even and odd frame equivalence is so *low* that the phonology threatens to drastically disrupt the abstract metrical framework. Two examples of such patterning are Wordsworth's "To Sleep (II)" and Rossetti's "Body's Beauty". The distribution of equivalence in these poems is compared with the forty English sonnets in Figures 4 and 5. Frames 1-3 in the Wordsworth sonnet, Frames 2-5 in the Rossetti sonnet and Frames 7-9 in *both* sonnets show exactly the *opposite* of the expected patterning. In the Rossetti poem in particular, the distribution of even and odd frame equivalence approximates to some extent the type of distribution we might expect to find in English *prose*. Particularly striking is the unusually high degree of equivalence in Frames 5, 7 and 9 of this poem. Frame 7 even exceeds the amount of equivalence in Frame 10, despite the fact that four of the poem's rhyme-pairs occur in the latter. One of the factors contributing to the obscuration of the expected contrast between even and odd frame equivalence in this poem is its monosyllabicity. Over 91 % of the words are monosyllables, while only 69.1 % of those in the sonnet by Keats are monosyllabic. In German, this degree of obscuration of contrast, with its attendant disruption of the meter, virtually never occurs in the nineteenth century, and even in the twentieth century it ap-

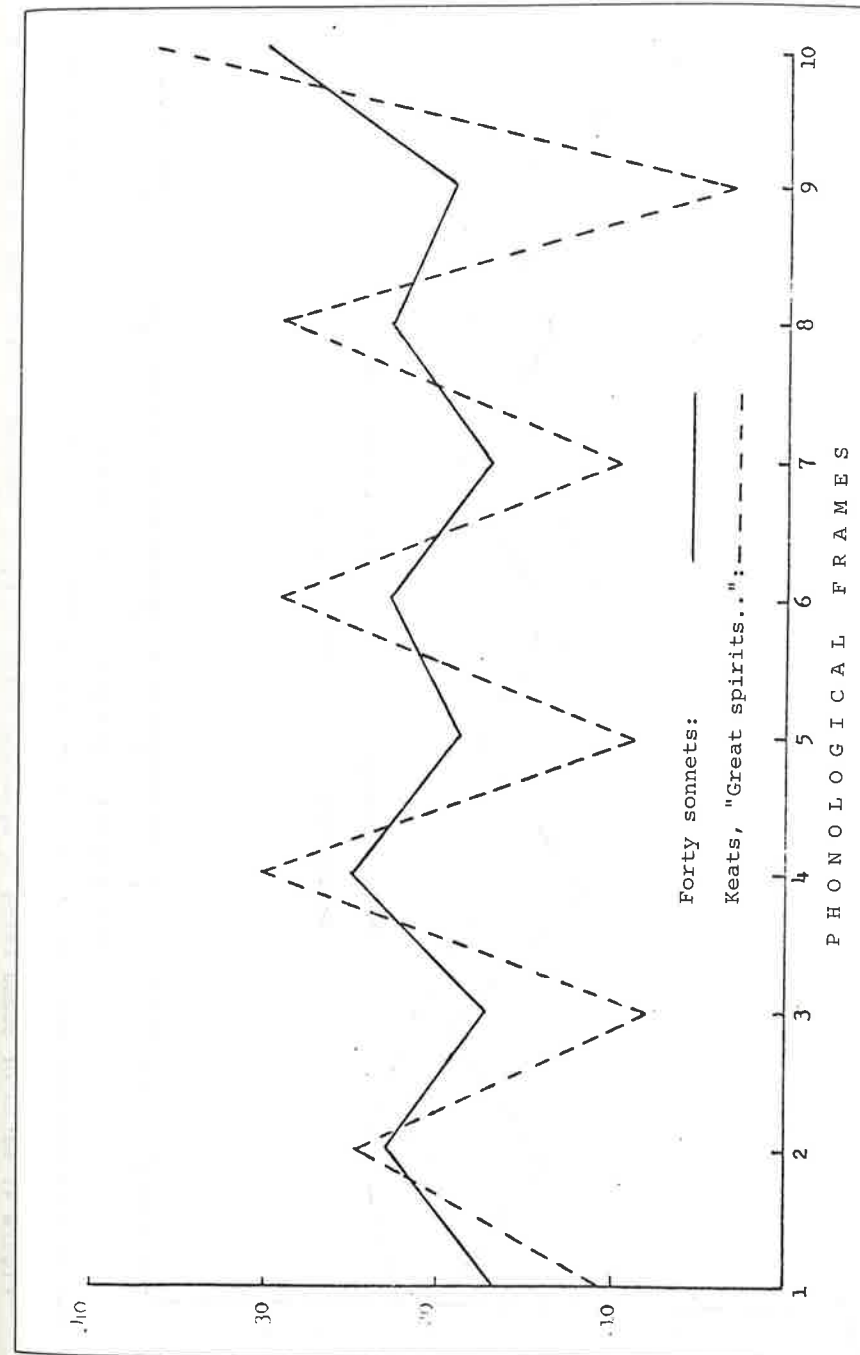


Figure 3: Ratio of Sound Pairs to Syllable Pairs in Forty English Sonnets and in Keats' "Great spirits now on earth ..."

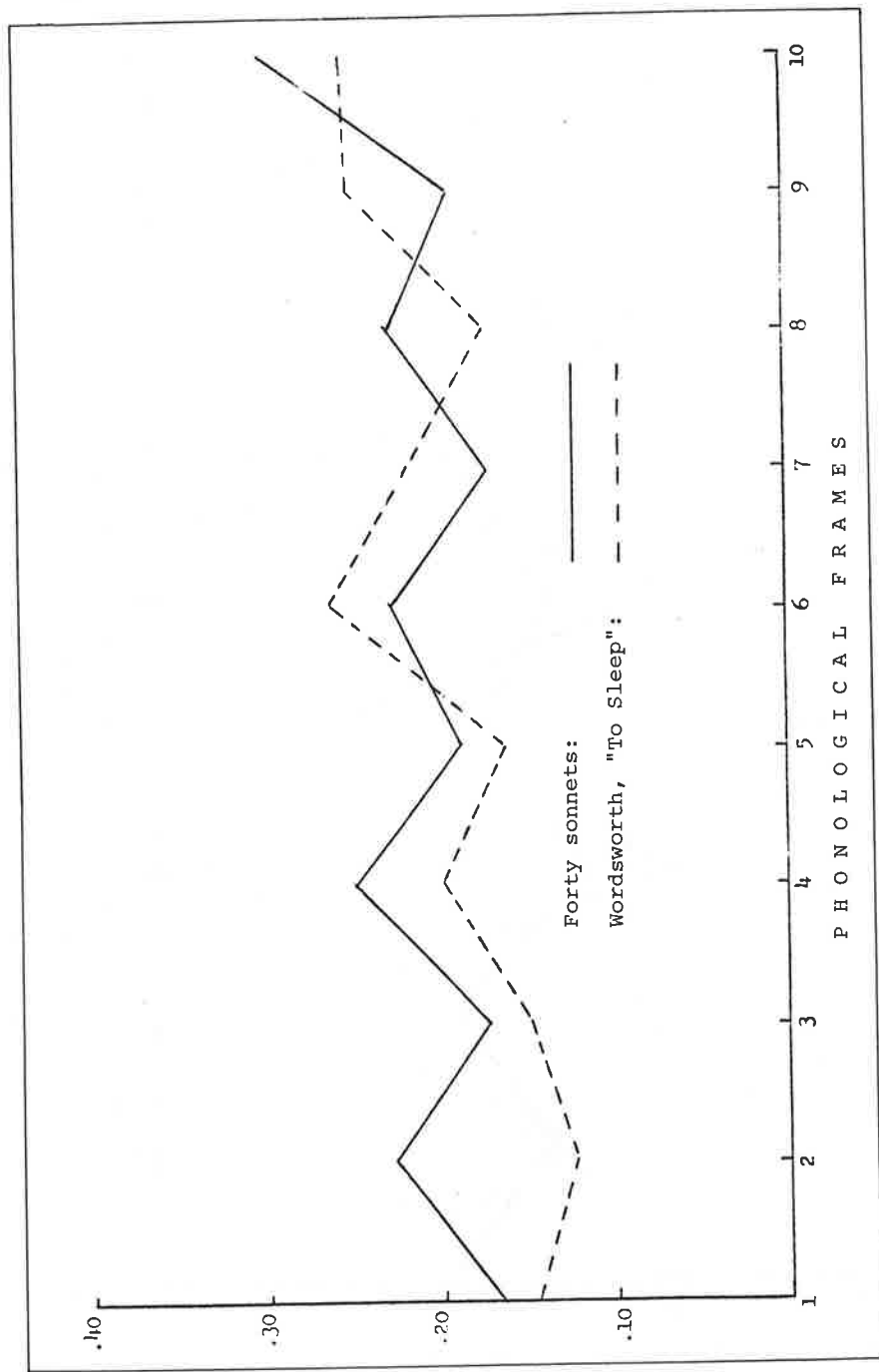


Figure 4: Ratio of Sound Pairs to Syllable Pairs in Forty English Sonnets and in Wordsworth's "To Sleep" (II)

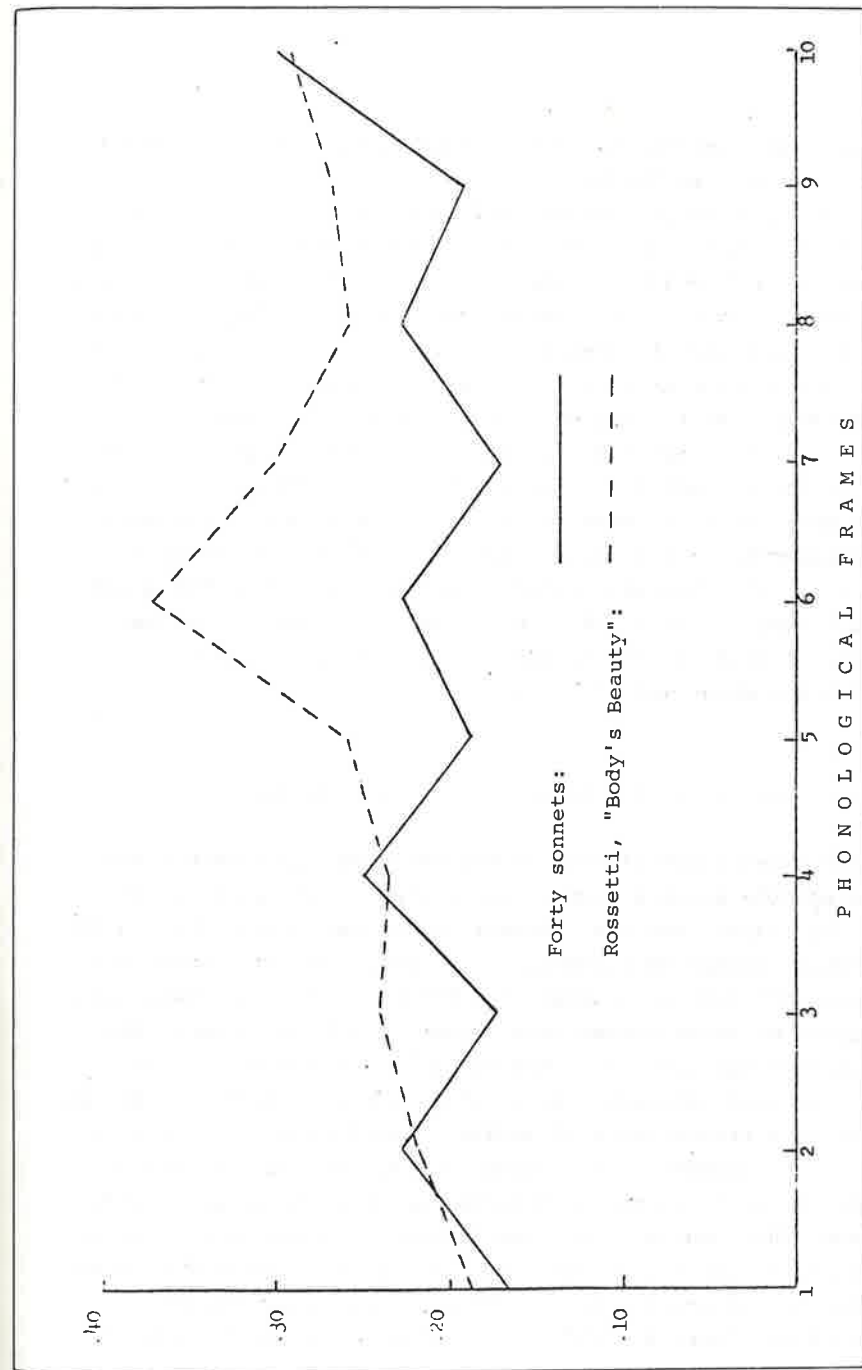


Figure 5: Ratio of Sound Pairs to Syllable Pairs in Forty English Sonnets and in Rossetti's "Body's Beauty"

pears to be limited to a very few poets such as Rainer Maria Rilke and Bertold Brecht.

Part C of Table 4 showed that the total equivalence in Frames 1-10 in the sonnets by the late nineteenth century poets Swinburne and Rossetti significantly exceeds that of the early nineteenth century poets Wordsworth and Keats. Table 5 provides a closer look at each poet by giving, in rank order, the coefficient of equivalence in Frames 1-10 for each of the forty sonnets. The ten segments from Swinburne's tragedy *Atalanta in Calydon* are included for comparison, and the poets are arranged chronologically from left to right according to the approximate dates of composition of the poems. While the sonnets of Wordsworth show a fairly wide range of values, those by Keats are all clustered together at the low end of the scale with a range of only .042. The sonnets of Rossetti and Swinburne, as well as the passage from *Atalanta in Calydon*, tend toward the upper end of the scale.

9. Vertical Patterning in German and English Verse

I have shown elsewhere that nineteenth century German verse consistently shows a higher equivalence coefficient in the primary frames (vertical patterning between lines) than in the secondary frames (horizontal patterning) even when end-rhyme is excluded from consideration (CHISHOLM 1976: 13). Subsequent analysis of other German poets shows that this tendency prevails into the twentieth century as well although the difference becomes markedly less in the sonnets of Rilke and Brecht. Since this predominance of vertical patterning throughout the verse line appears to be a general characteristic of German iambic verse since the eighteenth century, the question arises whether this feature might not be common to the iambic meters of other Germanic languages, such as English. To test this hypothesis, the primary and secondary frames were compared in each of the forty English sonnets. Table 6 shows the coefficient of equivalence in primary and secondary frames both with and without end-rhyme in both the German and English sonnets.

Table 5: Coefficients of Equivalence in Frames 1-10 for Four English Poets

	Wordsworth	Keats	Rossetti	Swinburne, Sonnets	Swinburne, Atalanta		Wordsworth	Keats	Rossetti	Swinburne, Sonnets	Swinburne, Atalanta
.28					.282 (IV)	.22	.217 (VIII) .216 (IV) .214 (V)	.219 (III) .217 (VIII) .214 (IV) .213 (VII) .210 (VI) .207 (X)	.213 (III) .208 (II) .208 (VII)	.213 (IX) .208 (V) .207 (VI)	.215 (VII) .213 (VI) .207 (X) .205 (III) .202 (V)
.27				.268 (III)		.21			.213 (III) .208 (II) .208 (VII)	.208 (V) .207 (VI)	
.26			.261 (IX) .257 (X)		.261 (I)	.20	.200 (VI)		.200 (I) .200 (VIII)	.200 (VII)	
.25				.252 (IV) .247 (X)		.19	.192 (III)	.194 (I) .189 (IX) .187 (V)	.193 (IV)		.190 (IX)
.24	.239 (IX) .236 (I)		.235 (VI)	.240 (II) .234 (VII) .232 (I)	.239 (VIII)	.18	.178 (II) .177 (X)	.177 (II)			
.23						.17					
.22			.222 (V)			.16	.158 (VII)				
						.15					

Table 6: Phonological Equivalence in Primary (with and without End-rhyme) and Secondary Frames in English and in German Sonnets

I. English	(Primary Frames 1-4)		(Secondary Frames 2,4,6,8)
	Slots 1-10(E)	Slots 1-8(-E)	Slots 1-10
Wordsworth	.262	.224	.215
Keats	.305	.248	.231
Rossetti	.292	.245	.236
Swinburne, <i>Sonnets</i>	.320	.265	.263
(Swinburne, <i>Atalanta</i>)	.249 (-E)		.261)
II. German			
Goethe, <i>Sonnets</i>	.37	.34	.30
Goethe, <i>Iphigenie</i>	.38 (-E)		.33
Schlegel	.42	.37	.32
Brentano	.42	.37	.31
Heine	.44	.44	.35
Mörike	.42	.36	.30
Meyer	.44	.41	.37
Keller	.46	.38	.36
Heym	.48	.44	.37
Rilke	.38	.36	.32
Brecht	.38	.33	.31

For each poet the number of syllable-pairs available for comparison in the primary frames (1-4) was 4600 for Slots 1-10 and 3680 for Slots 1-8; in the first four secondary even frames (2, 3, 6, 8) 5400 syllables were available for comparison. The table shows clearly that vertical patterning, other than that attributable to endrhyme, is *not* a dominant feature in the sonnets of Wordsworth, Keats, Rossetti and Swinburne. Whereas in German the differences between the coefficients of equivalence for primary frames in Slots 1-8 and the secondary frames are highly significant, the slight differences in the English sonnets are not statistically significant at the 5 % level of probability. Only when end-rhyme is included does English verse show significant increase in the primary frames.

10. Phonological Development in English and German Sonnets

A poem is perceived by the listener or reader not all in one instant, but rather as a progression of shifting relationships from beginning to end. This development can be analyzed on three levels: the verse line, the stanza and the whole poem. The phonology participates in this process, and for certain types of poetry the amount of phonological equivalence in a given frame or group of frames may differ significantly from one stage of the poem's development to another. Analysis of nineteenth century German sonnets, for example, showed that metrically disruptive odd frame equivalence decreases in the second half of sonnets by Brentano and Mörike, but *increase* in sonnets by Heine and C.F. Meyer. It also showed a considerable decrease in Frame 1 equivalence in the second half of the *line* in the sonnets by Goethe, and to a lesser extent Brentano. Table 7 gives the total coefficient of equivalence in both odd and even frames in the first and second half (lines 1-7 and 8-14 respectively) of the English sonnets by Wordsworth, Keats, Rossetti and Swinburne.

Whereas the Romantic poets Wordsworth and Keats both show a decrease in both odd and even frame equivalence in the se-

Table 7: Coefficient of Equivalence in Odd and Even Frames in the First and Second Half of Sonnets by Four English Poets

	First Half (lines 1-7)	Second Half (lines 8-14)	% Increase
1. Wordsworth			
Odd Frames :	605	536	- 11.40 %
<u>Even Frames</u> :	<u>822</u>	<u>660</u>	- 19.71 %
Total (Frames 1-10) :	1427	1196	- 16.19 %
2. Keats			
Odd Frames :	561	485	- 13.55 %
<u>Even Frames</u> :	<u>791</u>	<u>762</u>	- 3.67 %
Total (Frames 1-10) :	1352	1247	- 7.77 %
3. Rossetti			
Odd Frames :	585	673	+ 15.04 %
<u>Even Frames</u> :	<u>758</u>	<u>808</u>	+ 6.60 %
Total (Frames 1-10) :	1343	1481	+ 10.28 %
4. Swinburne			
Odd Frames :	656	780	+ 18.90 %
<u>Even Frames</u> :	<u>847</u>	<u>907</u>	+ 7.08 %
Total (Frames 1-10) :	1503	1687	+ 12.24 %

cond half of the sonnet, the Victorian poets Rossetti and Swinburne show a sizable *increase* in the second half, most of which is attributable to metrically disruptive odd frame equivalence. Since the "line" functions as an important rhythmic entity in verse, it is of interest to determine whether certain types of equivalence occur more or less frequently toward the end than at the beginning of the line. It was noted earlier that both poetry and prose tend to avoid Frame 1 equivalence, so that this frame usually shows the least equivalence of all frames in a given text. Table 8 gives the distribution of syllables containing Frame 1 equivalence from the beginning to the end of the line in the forty English sonnets under consideration. The sonnets of Wordsworth show a considerable *decrease* in the second half of the line, while those of Swinburne and Rossetti show an increase. Because of the small number of equivalences for each poet, only the results for Wordsworth are statistically significant; further work will be necessary to determine whether Rossetti and Swinburne consistently increase metrically disruptive Frame 1 equivalences in the second half of the line.

The following generalizations about English and German verse are supported by the results of this study:

- (1) Whereas vertical sound patterning predominates *throughout the verse line* in German iambic verse, the nineteenth century English sonnets show no statistically significant increase in vertical patterning other than that attributable to end-rhyme.
- (2) English verse, in part because of its greater monosyllabicity, shows substantially less contrast between even and odd frame equivalence than does German verse.
- (3) There is a marked tendency - due in part to stylistic factors - to avoid equivalence between contiguous syllables in both prose and poetry in English and German. Excessive Frame 1 equivalence would weaken an underlying alternation pattern which exists as a *tendency* in prose and as a conscious metrical device in verse.

Table 8: Distribution of Syllables Continuing Frame 1 Equivalence in the First and Second Half of the Verse Line

	Metrical Slots												
	First Half of Line						Second Half of Line						% Increase
	1	2	3	4	5	Total	6	7	8	9	10	Total	
	42	41	45	39	36	203	38	32	26	31	36	163	
Wordsworth	30	28	35	44	38	175	37	39	33	32	33	174	- 19.70 %
Keats	54	47	39	35	39	214	42	39	46	48	55	230	- 0.57 %
Swinburne	40	42	32	32	38	184	33	31	43	50	44	201	+ 7.48 %
Rossetti													+ 9.24 %

- (4) The English Victorian poets Swinburne and Rossetti are distinguished from the Romantic poets Wordsworth and Keats by (a) a significantly greater coefficient of equivalence in Frames 1-10 and (b) an *increase* in even and especially odd frame equivalence in the second half of the sonnet.
- (5) The results for both languages offer strong counterevidence to B.F. SKINNER's theory of "formal perseveration" in verse.

The usefulness of this computer-assisted approach to iambic verse in two languages suggests that its extension to other verse forms and other languages will not only reveal stylistic distinctions between poets and periods, but could also lead to the discovery of universal prosodic patterns which inhere in the poetic phonology of diverse languages.

Notes

- 1 An alternative solution would be to attach the linking-r to the beginning of the following word, so that the data would show only prevocalic /r/.
- 2 The significance of the difference can be checked by means of the chisquare test. For English, the difference between blank verse and the sonnets is highly significant at the 5 % level of probability ($\chi^2 = 12.40$), while in German the difference is not quite significant ($\chi^2 = 2.86$).
- 3 As early as the seventeenth century Martin OPITZ proscribed sequences such as "Die, die dir diese Dinge sagen." See M. OPITZ (1624).

References

Chisholm, David

- 1976 Phonological Patterning in German Verse. Computers and the Humanities 10, 5-20.

Gimson, A.G.

- 1962 An Introduction to the Pronunciation of English, London, Arnold.

Jones, Daniel

- 1956 An Outline of English Phonetics, 8th Edition, Cambridge, Heffer.

Opitz, Martin

- 1624 Buch von der deutschen Poeterey, Breslau, David Müller. (Reprinted in: Neudrucke deutscher Literaturwerke, Neue Folge 8, Tübingen, Niemeyer, 1963.)

Skinner, B.F.

- 1941 A Quantitative Estimate of Certain Types of Sound-Patterning in Poetry. American Journal of Psychology 54, 64-79.

TOWARDS A LAW OF WORD REPETITIONS IN
TEXT-BLOCKS

G. Altmann, V. Burdinski, Bochum

Abstract

The distribution of text-blocks containing x words (or other text units) is derived in the form of the negative hypergeometric (Beta-binomial) distribution. Tests are performed on English, Russian, Hungarian, Bulgarian and Indonesian data. Factors that are probably in relation to the parameters and the possible uses of this law are considered.

1. It is impossible to construct theories without laws, since according to the conception of all empirical sciences theories are systems of laws. Laws are hypothetical statement of general validity in a particular domain of reality, containing no observational concepts, derived from plausible assumptions which can be regarded as axioms, or which can be derived from a theory (even an embryonic one), i.e. can be derived from other laws, and can be sufficiently corroborated empirically (cp. Bunge 1961; 1967; 305-279). This rather casually adopted definition causes immediate rejection of empirical generalizations, systems of concepts and ad hoc hypotheses (e.g. rules) as potential laws, and this is probably the reason why in (text)linguistics (erroneously referred to as text theory) laws are at least for the time being still so rarely met with. Text is a complex entity in which extremely varied factors play a rôle (language, style, genre, dialect, age, status, etc.) so that it seems possible to set up only empirical generalizations. But if one considers that texts, independent of the language and personality of the author - represent a class of objects having common properties which come into existence by

virtue of the existence of such laws, one is encouraged to derive hypotheses about text construction from assumptions. At the moment it is only possible to set up and test "simple" hypotheses. In this article we shall attempt - without taking any factors into account at all - to derive a theoretical frequency distribution holding to all word distributions in text blocks, and to test it against data. We shall restrict ourselves to the distribution of a single word in a text block.

2. The first investigations along these lines were performed, as far as we have been able to discover, by Frumkina (1962). She examined the distribution of different words ("bez" without, "den'" day, "drugoј" another, "novyj" new, "očen'" very, "zdes'" here, "do" until, "meždu" between, "pri" at, "nad" above, "ili" or, "znat'" to know) in text-blocks of 1000 words in the work of Pushkin. Since the data are not easily accessible we present them in Table 1. Frumkina probably proceeded from the assumption that for entities with very low probability of occurrence "the law of small numbers" held and tried to fit the Poisson distributions to these data. The χ^2 test for goodness of fit and the respective probabilities are presented in Table 2. In three cases (meždu, nad, ili) small frequencies below 1.00 have been considered because of the small number of degrees of freedom. The results show that in 5 cases out of 12 the hypothesis of the Poisson distribution can be rejected. The fact that the mean and the variance of the Poisson distributions are equal should have been sufficient for the rejection of the Poisson distribution as inadequate (as a general model).

Brainerd (1972) was more successful in applying the Poisson distribution to the frequencies of the English article in 39 different samples of text-blocks consisting of 50 words respectively. He demonstrated that with one exception (a sample from J. Cheevers Wapshot Chronicle) all data fit the Poisson distribution (this holds for samples not tested by Brainerd as well). He decided to use the Poisson distribution because in all samples the mean and the variance yielded approximately the same value. The only exception was explained by Brainerd with

Table 1. The distribution of Russian words in passages of 1000 words according to Frumkina (1962)

Number of passages												
Number of words	bez without	den' day	drugoј another	novyj new	očen' very	zdes' here	do until	meždu between	pri at	nad above	ili or	znat' to know
0	36	36	28	88	59	70	44	104	76	97	101	33
1	42	26	34	28	26	25	35	5	24	11	8	27
2	23	16	28	11	13	11	17	1	9	2	1	14
3	6	17	15		10	2	6		1			12
4	3	11	4	7		0	6					7
5					4	2	1					4
6		4										
7												
8												
9												3
Σ	110	110	109	134	112	110	109	110	110	110	110	100
m ₁ ¹	1,0727	1,6091	1,3853	0,5821	0,9286	0,5727	1,0642	0,0636	0,4091	0,1364	0,0909	1,6150
m ₂	0,9947	2,5590	1,2460	1,0343	1,7181	0,9174	1,4179	0,0778	0,4599	0,1541	0,1008	3,1093

Table 2. Fitting the Poisson distribution to the data of Frumkina

Wort	χ^2	df	$P \approx$
bez	0,40	2	0,82
den'	21,41	3	<0,0005
drugoј	1,72	3	0,63
novyj	8,26	2	0,018
očen'	17,60	2	<0,0005
zdes'	4,67	1	0,03
do	3,43	2	0,18
meždu	1,95	1	0,17
pri	2,81	1	0,096
nad	1,88	1	0,18
ili	0,90	1	0,35
znat'	19,55	3	<0,0005

the assumption that the passages sampled were heterogeneous and represented perhaps a superposition of several Poisson distributions. By compounding three Poisson distributions he arrived at a good fit. It would probably be possible to achieve better fits to Frumkina's data using the compound Poisson distribution if one assumes for example that a particular word has several meanings which must be examined separately. This would, however, make semantic complications unavoidable.

Piotrowski (1975: 111-119) has shown that several Soviet authors tried to fit a posteriori the Poisson, the normal and the lognormal distribution and discussed the results. The data were not accessible to us. In this contribution we shall try to pursue another avenue of approach.

3. Let us consider any linguistic entity A (affix, word, phrase etc.) occurring in a language generally with probabi-

lity p and let us denote with n the number of positions in a passage in which A could occur. With context-free occurrence of A the p is constant, with context-dependent occurrence of A we have in most places of the passage $p = 0$, in other positions, depending on context, $p > 0$ but not necessarily constant. If p were constant then the probability to find A exactly x times in these n positions would obviously be

$$f(x|p) = \binom{n}{x} p^x (1-p)^{n-x} \quad x = 0, 1, \dots, n \quad (1)$$

representing the binomial probability distribution. Since p is different in each of the potential n positions the resulting probability to find exactly x As depends on p , i.e. it is a conditional probability. The parameter p can be considered as a continuous random variable whose domain lies in the interval $<0,1>$. In order to find the common distribution of x and p we take

$$f(x, p) = f(x|p)f(p)$$

from which the marginal distribution of x can be computed as

$$f(x) = \int_0^1 f(x|p)f(p)dp. \quad (2)$$

The distribution of p is not known. It cannot be derived theoretically because there is no such thing as the population of all texts of a language: there are merely sublanguages (a term introduced by Soviet authors) having their own (perhaps very different) distributions of p . Therefore in order to keep the theoretical background flexible we choose the Beta distribution as being general enough and allowing different text factors to be brought into relation to its parameters. We write the Beta distribution in the form

$$f(p) = \frac{1}{B(M, K-M)} p^{M-1} (1-p)^{K-M-1}, \quad 0 < p < 1 \quad (3)$$

where

$$B(a,b) = \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)} \quad \text{and} \quad \Gamma(k) = (k-1)!$$

Substituting (1) and (3) into (2) and integrating we obtain

$$\begin{aligned} P(X=x) = f(x) &= \int_0^1 \binom{n}{x} p^x (1-p)^{n-x} \frac{1}{B(M, K-M)} p^{M-1} (1-p)^{K-M-1} dp = \\ &= \binom{n}{x} \frac{1}{B(M, K-M)} \int_0^1 p^{M+x-1} (1-p)^{K-M+n-x-1} dp = \\ &= \binom{n}{x} \frac{B(M+x, K-M+n-x)}{B(M, K-M)}. \end{aligned}$$

Writing the Beta expressions as above and arranging them in form of binomial coefficients we obtain

$$\begin{aligned} P(X=x) &= \frac{\binom{M+x-1}{x} \binom{K-M+n-x-1}{n-x}}{\binom{K+n-1}{n}} = \\ &= \frac{\binom{-M}{x} \binom{-K+M}{n-x}}{\binom{-K}{n}}, \quad x = 0, 1, \dots, n \\ &\quad 0 < M < K \end{aligned} \quad (4)$$

The probability of finding exactly x times A in a text-block equals (4) and the expected number of text blocks containing exactly x entities A is given as $NP(X=x)$, where N is the number of text blocks sampled. Distribution (4) is called the negative hypergeometric (NH) (or Beta-binomial) distribution and represents a special case of the generalized hypergeometric distribution (Type IIa according to Kemp & Kemp 1956a; cf. also Johnson & Kotz 1969; Ord 1972). The parameters M and K are positive, not necessarily integers, $K > M$.

The NH distribution has three important limiting cases:

- (i) If $K \rightarrow \infty$, $M \rightarrow \infty$ so that $M/K \rightarrow p$ then it tends to the binomial distribution as given in (1).
- (ii) If $K \rightarrow \infty$, $n \rightarrow \infty$ so that $K/(K+n) \rightarrow p$ then it tends to the negative binomial distribution

$$P(X=x) = \binom{M+x-1}{x} p^M (1-p)^x, \quad x = 0, 1, \dots$$

(iii) If $M \rightarrow \infty$, $K \rightarrow \infty$, $n \rightarrow \infty$ so that $Mn/K \rightarrow \lambda$ then it tends to the Poisson distribution

$$P(X=x) = \frac{e^{-\lambda} \lambda^x}{x!}, \quad x = 0, 1, \dots$$

In this way we obtain the original binomial distribution, the Poisson distribution, which was relatively well corroborated on English articles, and the negative binomial distribution as limiting cases. If we succeed in showing empirically that the NH distribution or its limiting cases agree with the data then we would have a candidate for the law of word distributions in text passages.

The following circumstances are to be considered here:

(a) The Beta distribution is a very general density function, the use of which in this case would be difficult to justify theoretically. The fact that it can be derived as a special case of the Pearson curves or as a function of the normal or the Gamma distributions could be of some help. It may also turn out that one can make do with one or several special cases of the NH distribution.

(b) The parameters of the distribution should also be derived theoretically from other factors of language or text or from the properties of the entity examined or from other laws. Piotrowski (1975) mentions on the basis of examinations of several Soviet authors (Maškina 1968; Bektaev & Lukjanenkov 1971; Paškovskij & Srebrjanskaja 1971) the following possible factors: (i) the relative frequency of a word in a frequency dictionary; (ii) the length of the passage in which the word distribution is examined; (iii) the word class and the lexical-grammatical character of the word; (iv) the relevance of the word for the theme.

Standard linguistics does not provide any help along these

lines so that for the time being we are restricted to the collecting and testing of data.

If there is a priori no possibility of determining the parameters, one must estimate them from the data. One could take n as the greatest number of occurrence of the examined entity in passages of the given length in the given language and thereby spare one degree of freedom with the test of fit. If this is not possible (preliminarily) then we set

$$\hat{n} + 1 \geq \text{number of classes of the random variable.}$$

If $f_0 \neq 0$ then

$$\hat{n} \geq x_{\max};$$

if $f_0 = 0, f_1 \neq 0$ then

$$\hat{n} \geq x_{\max} - 1 \text{ etc.}$$

The other parameters can be estimated by means of the moments

$$\mu'_1 = \frac{nM}{K}$$

$$\mu'_2 = \frac{nm(k-M)(K+n)}{K^2(K+1)}$$

as

$$\hat{K} = \frac{n(n\bar{x} - \bar{x}^2 - s^2)}{n(s^2 - \bar{x}) + \bar{x}^2}$$

and

$$\hat{M} = \frac{\hat{K}\bar{x}}{n}.$$

The better but much more lengthy maximum-likelihood method can be found in Kemp & Kemp (1956b). A chi-square test for fitting the NH distribution can be performed only if there are at least 5 classes since we must subtract 4 degrees of freedom.

The fitting of the NH distribution to the data of Frumkina is presented in Table 3. This table contains the expected frequencies $NP(X=x)$, where $N = \sum f_i$ = sample size and $P(X=x)$ was computed from (4) by means of the recursion formula

Table 3. The NH distribution and the test of fit for Frumkina's data

Expected number of passages NP (X=x)												
Number of words in passage	bez	den'	drugo j	novy j	očen'	zdes'	do	mezdu	pri	nad	ili	znat'
0	36,21	36,09	27,50	90,09	60,81	72,11	46,19	104	75,55	97	101	35,85
1	41,59	25,64	35,17	23,45	23,47	21,74	30,25	5	25,36	11	8	22,55
2	22,58	18,82	27,47	11,23	13,09	9,72	18,05	1	7,65	2	1	15,50
3	7,58	13,39	14,54	5,71	7,63	4,34	9,45		1,45			10,64
4	1,73	8,87	4,32	2,64	4,25	1,67	3,99					7,06
5	0,28	5,11		0,88	2,06	0,43	1,06					4,40
6	0,03	2,09			0,70							
n	9	6	4	5	6	5	5	2	3	2	2	8
K	150,62	3,2621	6,9788	2,9565	3,2050	3,9438	4,7752	2,8131	5,6295	3,6948	5,1766	3,9568
M	17,95	0,8748	2,4170	0,3442	0,4960	0,4517	1,0163	0,0895	0,7677	0,2519	0,2353	0,7988
χ^2	0,05	1,47	0,10	*	2,35	1,98	2,92					1,66
df	1	2	1		1	1	1	*	*	*	*	2
$P \approx$	0,83	0,48	0,76		0,13	0,17	0,09					0,44

$$P(X=0) = \frac{(K-M)(K-M+1) \dots (K-M+n-1)}{K(K+1) \dots (K+n-1)}$$

$$P(X=x+1) = \frac{(M+x)(n-x)}{(x+1)(K-M+n-1-x)} \cdot P(X=x)$$

With bez the highest class is to be understood as $NP(X \geq 6) = 0.03$, with znat' as $NP(X \geq 5) = 4.40$. The estimated parameters, the computed χ^2 values, the degrees of freedom (df) and the corresponding probabilities (P) are in the rows under the table. In samples marked with* (meždu, pri, nad, ili) the χ^2 test could not be performed because of small df but it is evident that the fit is good. Several fits could be further improved if a greater n were chosen or if other estimations were performed. It was, however, not our aim to find the best fit but to test the adequacy of the model.

Since Brainerds data with one exception follow the Poisson distribution it is sufficient to examine this exception. The observed and the computed frequencies are given in Table 4. As may be seen, the fitting of the NH distribution is good even in this case, although the empirical distribution is somewhat "pathological".

5. For the sake of the further testing of the model data from 3 other languages have been collected. Our hypothesis can be considered as corroborated by any empirical case if the NH distribution or some of its limiting distributions can be adequately fitted. Since the fitting of the limiting distributions is easier, one can use them if there are favourable circumstances. When choosing among them one can use the following criteria:

if $\bar{x} > s^2$ binomial

$\bar{x} = s^2$ Poisson

$\bar{x} < s^2$ negative binomial

where $\bar{x} = m_1' = \frac{1}{N} \sum x_i f_i$ and $s^2 = m_2' = \frac{1}{N} \sum (x_i - \bar{x})^2 f_i$ (we used here

Table 4. The distribution of the article in a sample of Cheevers Wapshot Chronicle (WC 1) according Brainerd (1972)

Number of article in passage	Number of passages		
	observed	Poisson	NH
0	8	1,54	7,62
1	8	6,14	9,89
2	12	12,23	10,71
3	11	16,26	10,72
4	14	16,22	10,20
5	4	12,93	9,29
6	8	8,60	8,07
7	8	4,90	6,62
8	6	2,44	5,01
9	2	1,08	3,30
10	2	0,66	1,58
$\lambda = 3,988$			n=10 K=3,6401 M=1,4517
χ^2			18,30
df			6
P $\approx$			0,0056
			5,61 6 0,47

the biased estimation of σ^2). The estimation of the parameters can be performed e.g. as follows:

$$\text{Binomial: } \hat{n} = x_{\max}$$

$$\hat{p} = 1 - \sqrt{\frac{\hat{n}}{f_0/N}}$$

where f_0/N is the relative frequency of the zero class.

Poisson: $\hat{\lambda} = \bar{x}$

Negative binomial: (1) $\hat{p} = \frac{\bar{x}}{s^2} - 2$; $\hat{k} = \frac{\bar{x}^2}{s^2 - \bar{x}}$

or iteratively from (2) $\frac{\bar{x}}{-\log \frac{f_0}{N}} = \frac{\hat{p}}{\log(1+\hat{p})}$; $k = \frac{\bar{x}}{\hat{p}}$

where $\hat{p} = \frac{1}{1 + \hat{p}}$.

Let us now turn to the use of the limiting distributions in a concrete case. In Table 5 the empirical distribution of the German article das is presented in 200 passages from "Deutschstunde" by S. Lenz (column 1). In columns 2 to 6 n will be successively increased and it can be seen that the fit improves. Since $\bar{x} < s^2$, we assume that with increasing n and K the negative binomial will yield the best fit. The estimation of p according to (1) gives $\hat{p} = 0.5663$ and by means of (2) $\hat{p} = 0.6254$. Computing the criteria $K/(K+n)$ for columns 2 to 6 one sees that it slowly tends to $\hat{p}$:

Column	2	3	4	5	6
$K/(K+n)$	0.4275	0.4449	0.4969	0.5364	0.5498

In column 7 the negative binomial distribution is presented. As can be seen, this fit is the best one.

For the German samples in Tables 6 and 7 we used only the negative binomial and in one case, namely with "bis" ("until" with temporary meaning) the binomial because of $\bar{x} > s^2$. The use of the χ^2 -test was not possible in each case but the agreement is even optically very good.

The theoretical frequency of the highest class $x_{\max}$ of the negative binomial distribution was computed as $\sum_{x=x_{\max}}^{\infty} NP_x^x$.

Table 5. The distribution of the article "das" in the nominative in passages from Lenz

Number of "das" in the passage	observed number of passages	NH						Negative binomial
		102,64	102,04	100,45	99,28	99,13	95,00	
0	95	47,49	48,38	50,72	52,28	52,64	56,44	
1	58	25,23	25,34	25,60	25,83	25,78	27,34	
2	28	13,32	13,11	12,60	12,33	12,20	12,24	
3	11	6,66	6,46	5,98	5,72	5,64	5,26	
4	3	3,04	2,94	2,72	2,59	2,57	2,20	
5	2	1,19	1,20	1,18	1,14	1,15	0,90	
6	1	0,37	0,41	0,48	0,49	0,51	0,37	
7	1	0,07	0,12	0,27	0,34	0,38	0,25	
≥ 8	1							
$\hat{n}$		8	9	15	30	60	$\hat{p} = 0,6254$	
$\hat{k}$		5,9740	7,2121	14,8159	34,7090	73,2888	$\hat{k} = 1,5862$	
$\hat{M}$		0,7094	0,7613	0,9383	1,0991	1,1604		
χ^2		5,64	4,89	3,30	2,47	2,30	0,29	
df		2	2	2	2	2	2	
$P \approx$		0,06	0,09	0,19	0,29	0,32	0,87	

$N = 200$; $m_1' = 0,9500$; $m_2 = 1,6775$.

Table 6. Distribution of "das" in passages from Lenz

number in the passage	article in the accusative	NB	demonstrative pronoun in the nominative	NB	demonstrative pronoun in the accusative	NB	relative pronoun in the nominative	NB	relative pronoun in the accusative	NB
0	46	46,00	85	86,67	122	122,40	150	149,88	170	170,01
1	57	52,32	66	61,50	51	49,08	42	42,32	27	26,43
2	31	40,62	28	30,64	15	18,24	7	6,86	2	3,18
3	31	26,64	11	13,10	9	6,60	1	0,94	1	0,38
4	16	15,87	8	5,14	2	2,36				
5	7	8,88	0	1,91	0	0,84				
6	7	4,76	2	1,04	1	0,48				
7	4	2,46								
≥ 8	1	2,45								
Σ	200	200	200	200	200	200	200	200	200	200
m ₁	1,9450		0,9950		0,6100		0,2950		0,1700	
m ₂	3,2320		1,3950		0,9279		0,3080		0,1911	
p̂		0,5848		0,7133		0,6574		0,9578		0,9146
k̂		2,7394		2,4751		1,1705		6,6923		1,8201
χ ²		4,87		1,38		0,94		*		*
df		5		2		1				
p ≈		0,43		0,50		0,34				

Table 7. Distributions of "in" and "bis" in passages from Lenz

number in the passage	Number of passages					
	"in" + dative	NB	"in" + accusative	NB	"bis" spatial	"bis" temporal
0	10	11,28	52	52,00	160	154
1	34	30,40	61	62,37	30	42
2	42	42,82	46	43,98	10	4
3	37	41,96	25	23,77		
4	38	32,10	8	10,89		
5	17	20,44	4	4,45		
6	12	11,25	4	2,54		
7	6	5,52				
8	2	2,44				
9	2	1,79				
m ₁	200	200	200	200	200	200
m ₂	3,0700		1,5200		0,2500	0,2500
	3,4951		1,8996		0,2875	0,2275
p̂		0,8784		0,7891		0,8000
k̂		22,1710		5,6865		1,0000
χ ²		2,94		1,10		*
df		6		3		
p ≈		0,81		0,78		*
						0,1225 n̂ = 2

Table 8. The distribution of Bulgarian words in text-blocks

number in the pass.	tova (demonstrative pronoun)				koeto (relative pronoun)				-to (definite article)			
	AM	NH	St	NH	AM	NH	St	NH	AM	NH	St	NH
0	43	44,97	79	79,24	58	58,99	84	84,12	1	2,37	9	9,65
1	26	24,88	17	16,17	24	21,90	14	13,76	8	5,31	16	16,46
2	19	14,25	3	3,95	11	11,33	2	2,12	7	8,00	18	18,86
3	5	7,97	1	0,64	4	5,65			8	10,06	20	17,70
4	3	4,25			3	2,12			9	11,29	17	14,43
5	1	2,13							13	11,69	11	10,39
6	2	0,97							17	11,32	4	6,59
7	1	0,58							13	10,32	1	3,62
8									5	8,88	2	1,64
9									8	7,19	2	0,56
10									3	5,44	1	0,11
11									3	3,79		
12									2	2,38		
13									1	1,28		
14									0	0,54		
15									2	0,14		
Σ	100	100	100	100	100	100	100	100	100	100	100	100
m' 1	1,15				0,70		0,18		5,28		3,01	
m' 2	2,09		0,31		1,03		0,19		9,77		4,11	
$\hat{n}$		10		3		4		2		15		10
$\hat{R}$		7,5426		5,5478		2,8287		5,2519		7,0323		8,4396
$\hat{M}$		0,8674		0,4808		0,4950		0,4727		2,7285		2,5403
χ^2		2,94				1,07				7,81		2,05
df		1		*		1		*		7		4
$p \approx$		0,09				0,30				0,35		0,73

St - Stanev, Antichrist; AM - Andrejčin, Mirčev, Slavističen Sbornik

Table 8 (cont.). The distribution of Bulgarian words in text-blocks

number in the pass.	do "up to"				v "in"			
	AM	NH	St	NH	AM	NH	St	NH
0	77	77,12	81	81,00	2	0,90	2	2,72
1	14	13,76	17	17,00	0	2,58	9	6,51
2	9	9,12	2	2,00	5	4,73	14	9,92
3					6	6,99	11	12,15
4					9	9,06	18	12,98
5					13	10,65	15	12,56
6					14	11,59	9	11,23
7					14	11,76	10	9,39
8					10	11,18	3	7,39
9					7	9,90	4	5,48
10					7	8,11	4	3,83
11					4	6,01	0	2,53
12					4	3,88	1	3,31
13					4	2,01		
14					1	0,66		
15								
Σ	100	100	100	100	100	100	100	100
m' 1	0,32		0,21		8,82		5,28	
m' 2	0,40		0,2059		9,21		9,19	
$\hat{n}$		2		2		14		20
$\hat{R}$		1,0488		2,4708		6,9600		12,9209
$\hat{M}$		0,1678		0,9944		3,3905		3,4111
χ^2		*		*		4,48		10,29
df						7		6
$p \approx$						0,72		0,11

Table 9. The distribution of Indonesian words in text-blocks

number in the pass.	dan "and"			dengan "with"		
	H	NH	A	H	NH	A
0	2	2,11	3	8	10,09	13
1	3	4,98	2	20	18,42	13
2	5	7,62	7	24	21,21	35
3	13	8,87	9	17	19,12	31
4	16	11,40	15	16	14,36	29,56
5	16	12,10	13	7	9,13	22,42
6	9	11,97	21	5	4,85	13,32
7	6	11,09	12	2	2,07	6,03
8	10	9,60	9	0	0,64	1,89
9	8	7,70	4	1	0,11	0,32
10	3	5,61	1			
11	6	3,58	1			
12	2	1,85	1			
13	1	0,61	1			
14			1			
Σ	100	100	100	100	100	100
m'_1	5,77		5,46	2,71		2,81
m'_2	8,56		6,69	3,25		2,70
$\bar{n}$		13			9	
$\bar{x}$		6,1964			10,1739	
$\bar{m}$		2,7503			3,0635	
χ^2		11,34			1,87	
df		7			3	
$P_{\approx}$		0,13			0,60	
						8
						13,5504
						4,7596
						6,36
						3
						0,096

H - Hartwardojo, Orang Buangan: A - Adiwinarto, Bintangpun tiada

In Table 8 some distributions from Bulgarian are presented. The differences between prose (E. Stanev, Antichrist, Sofia 1973) and scientific literature (L. Andrejčin & Mirčev (Eds.), Slavističen sbornik, Sofia 1973) can be observed in the differences of the parameters but the model itself holds true.

In Table 9 the Indonesian word "dan" (and) and "dengan" (with) in two novels (B. Adiwinarto, Bintangpun tiada, Bandung 1971 and H.S. Hartwardojo, Orang Buangan, Djakarta 1971) are examined.

6. The circumstances that word distributions in text-blocks are governed by a stochastic law shows that in spite of (stilistic and other) differences between text and in spite of rules which are responsible for the wellformedness of sentences, the construction of texts is governed by laws. In order to elaborate these laws the following procedures must be performed:

(a) In absence of a theory models for individual text phenomena must be derived from assumptions. At the beginning one will be forced to compensate for the lack of rigorous derivations at some stages by approximations or trial-and-error assumptions. In the course of the development they should be, of course, replaced with plausible, interpretable steps.

(b) The estimations of the parameters extracted at the beginning from the data must be later on replaced by derived quantities which are related to other properties of language or text.

(c) The lawlike statements should be interrelated, i.e. one must construct a system of text laws that can be later on derived from an axiomatic background. We find ourselves here on the threshold of a "physics of text", which is developing into a very extensive discipline (cf. e.g. Orlov & Boroda & Nadarejšvili 1980).

A satisfactorily elaborated theory of word repetitions in text-blocks could be of help for the following procedures (cf. Piotrowski 1975):

- (i) mechanical identification of the membership of a word to a word class;

- (ii) identification of terminologically and semantically dominating units;
- (iii) ascertaining of the stylistic relevance of the word;
- (iv) measurement of the stylistic individuality of a text;
- (v) diagnosing the foci of some psychic diseases;
- (vi) construction of learning automata.

Literatur

- Bektaev, K.B., Luk'janenkov, K.F., O zakonach raspredelenija edinic pis'mennoj reči. In: Piotrowski, R.G. (Ed.), Statistika reči i automatičeskij analiz teksta. Leningrad, Nauka 1971, 47-112.
- Brainerd, B., Article Use as an Indicator of Style among English-Language Authors. In: Jäger, S. (Hrsg.), Linguistik und Statistik. Braunschweig, Vieweg 1972, 11-32.
- Bunge, M., Kinds and Criteria of Scientific Laws. Philosophy of Science 28, 1961, 260-281.
- Bunge, M., Scientific Research I. Berlin, Springer 1967.
- Frumkina, R.M., O zakonach raspredelenija slov i klassov slov. In: Mološnaja, T.N. (Hrsg.), Strukturno-tipologičeskie issledovanija. Moskva, AN SSSR, 1962, 124-133.
- Johnson, N.L., Kotz, S., Discrete Distributions. Boston, Houghton Mifflin 1969.
- Kemp, C.D., Kemp, A.W., Generalised Hypergeometric Distributions. Journal of the Royal Statistical Society, Series B 18, 1956a, 202-211.
- Kemp, C.D., Kemp, A.W., The Analysis of Point Quadrat Data. Australian Journal of Botany 4, 1956b, 167-174.
- Maškina, L.E., O statističeskich metodach issledovanija leksiko-grammatičeskoj distribucii. Minsk 1968 (Diss.).
- Ord, J.K., Families of Frequency Distribution. London, Griffin 1972.

- Orlov, In.K., Boroda, M.G., Nadarejšvili, I.Š., Sprache, Text, Kunst. Bochum, Brockmeyer 1980 (in press).
- Paškovskij, V.E., Srebrjanskaja, I.I., Statističeskie ocenki pis'mennoj reči bol'nych šizofreniej. In: Inženernaja lingvistika. Leningrad 1971.
- Piotrowski, K.G., Tekst, mašina, čelovek. Leningrad, Nauka 1975.

KANONISCHE ANALYSE VON KONTINGENZTAFELN¹

J. Lauber, J. Linhart, Prag

In diesem Aufsatz möchten wir den Leser mit einer der weniger bekannten Analysetechniken von Kontingenztafeln bekannt machen. Formal geht es um ein Verfahren, das mittels einer Transformation der Klassifikationsmerkmale (bzw. Dimensionen) der Kontingenztafel in eine zweidimensional normalverteilte Zufallsvariable den Korrelationskoeffizienten zwischen den Merkmalen maximiert und das für jede Kategorie (Ausprägungsstufe) der beiden Merkmale einen Score bestimmt.

Die Technik kann im Hinblick auf das Ziel und den verwendeten mathematischen Apparat unter die multivariaten Methoden der Analyse von soziologischen Daten eingeordnet werden. Die Vorgehensweise wurde sehr allgemein formuliert, so daß sie für alle Typen von Kontingenztafeln verwendet werden kann - für Tabellen, die durch Klassifizieren aufgrund von ordinal- und intervallskalierten Merkmale entstanden sind - wobei die Tabellen beliebig groß sein können. Die Interpretation der Ergebnisse ist nicht immer einfach, sie erfordert das Verstehen der Grundlagen der Methode. Die vorgestellte Technik erlaubt uns vor allem einen kompakten Überblick über die Struktur der Zusammenhänge zwischen den Merkmalen der Kontingenztafel. Dieser Überblick sollte durch eine detailliertere Analyse vervollständigt werden.

MATHEMATISCHE GRUNDLAGE DER KANONISCHEN ANALYSE

Die Methode basiert vor allem auf dem Lancaster-Satz (LANCASTER 1957; KENDALL, STUART 1973):

Wir wollen annehmen, daß die Zufallsvariable $F = (X, Y)$ eine zweidimensionale Normalverteilung mit der Kovarianzmatrix

$$\Sigma = \begin{bmatrix} \sigma_1^2 & \rho\sigma_1\sigma_2 \\ \rho\sigma_2\sigma_1 & \sigma_2^2 \end{bmatrix},$$

den Varianzen $D(X) = \sigma_1^2$ und $D(Y) = \sigma_2^2$ und der Kovarianz, bzw. der Korrelation

$$\text{cov}(XY) = \rho\sigma_1\sigma_2 \text{ und } \text{corr}(XY) = \rho$$

darstellt.

Transformiert man die einzelnen Komponenten der gegebenen zweidimensionalen Verteilung F , d.h.

$$X' = X'(X) \text{ und } Y' = Y'(Y),$$

dann gilt unter sehr allgemeinen Bedingungen (Endlichkeit der Streuungen der Größen X' und Y')

$$\text{corr}(X'Y') \leq |\rho|,$$

d.h. die Korrelation der transformierten Verteilung kann im absoluten Wert die Korrelation der ursprünglichen Normalverteilung nicht überschreiten.

Der Lancaster-Satz gilt auch für nicht kategorisierte Größen. Betrachten wir weiter eine Kontingenztafel (Tab. 1), die die absoluten Vorkommenshäufigkeiten aller Paare von Kategorien (Ausprägungsstufen) der Merkmale (Dimensionen) X und Y zusammen mit Zeilen- und Spaltensummen enthält (Dimension X soll r Klassen haben, und Dimension Y soll c Klassen haben).

Gegenstand der weiteren Überlegungen ist die Frage, auf welche Art und Weise den einzelnen Ausprägungen der Dimensionen X und Y bestimmte Zahlenwerte (Scores) zugeschrieben werden sollen, damit der Korrelationskoeffizient bei gegebenen Häufigkeiten sein Maximum erreicht.

Tabelle 1.

Merkmal Y					
Klasse	1	2	...	c	Σ
1	n_{11}	n_{12}	...	n_{1c}	$n_{1.}$
2	n_{21}	n_{22}	...	n_{2c}	$n_{2.}$
.	.	.	...	.	.
.	.	.	...	.	.
X	.	.	...	.	.
r	n_{r1}	n_{r2}	...	n_{rc}	$n_{r.}$
Σ	$n_{.1}$	$n_{.2}$	...	$n_{.c}$	n

Mit anderen Worten: wir wollen die Werte der diskreten Zufallsgrößen X und Y (zwecks Eindeutigkeit standardisiert), d.h. Zahlen (Score) x_i und y_j ($i = 1, 2, \dots, r$) und ($j = 1, 2, \dots, c$) so finden, daß die Randverteilung

$$\begin{array}{l}
 \text{X:} \quad \begin{array}{c} x_i \\ \hline P(X = x_i) \end{array} \parallel \begin{array}{ccc} x_1 & x_2 & \dots & x_r \\ \hline \frac{n_{1.}}{n} & \frac{n_{2.}}{n} & & \frac{n_{r.}}{n} \end{array} \\
 \\
 \text{Y:} \quad \begin{array}{c} y_j \\ \hline P(Y = y_j) \end{array} \parallel \begin{array}{ccc} y_1 & y_2 & \dots & y_c \\ \hline \frac{n_{.1}}{n} & \frac{n_{.2}}{n} & & \frac{n_{.c}}{n} \end{array}
 \end{array}$$

die Einheitsvarianzen

$$D(X) = \sum_{i=1}^r \frac{n_{i.}}{n} x_i^2 = 1,$$

$$D(Y) = \sum_{j=1}^c \frac{n_{.j}}{n} y_j^2 = 1 \text{ hat,}$$

wobei der Korrelationskoeffizient

$$\text{corr}(X, Y) = \sum_{i=1}^r \sum_{j=1}^c \frac{n_{ij}}{n} x_i y_j \longrightarrow \text{maximal ist.}$$

Eine unmittelbare Ausnutzung des angeführten Lancaster-Satzes, d.h. eine direkte Konstruktion der zweidimensionalen Normalverteilung, kommt nicht in Frage. Das Auffinden der optimalen Werte x_i und y_j (hinsichtlich der Maximalisierung des Korrelationskoeffizienten) können wir als eine Aufgabe der Suche nach einem gebundenen Extremum unter Randbedingungen lösen.

Bezeichnet man

$$f(x_1, x_2, \dots, x_r, y_1, y_2, \dots, y_c) = \sum_i \sum_j \frac{n_{ij}}{n} x_i y_j = \text{corr}(XY),$$

$$g_1(x_1, x_2, \dots, x_r) = \sum_i \frac{n_{i.}}{n} x_i^2 - 1 = 0,$$

$$g_2(y_1, y_2, \dots, y_c) = \sum_j \frac{n_{.j}}{n} y_j^2 - 1 = 0,$$

so bekommt man die Aufgabe, das Maximum von f unter den Bedingungen $g_1 = g_2 = 0$ zu suchen. Diese Aufgabe kann man mit Hilfe von Lagrange-Multiplikatoren lösen. Da die Kenntnis der Art der Bedingungen und Beziehungen, die von x_i und y_j erfüllt werden, zur Interpretation wichtig ist, führen wir die zur Lösung der Extremwertaufgabe verwendeten Hauptverfahren auf.

Zunächst bilden wir die Lagrange-Funktion:

$$\begin{aligned}
 \phi = & f(x_1, x_2, \dots, x_r, y_1, y_2, \dots, y_c) + \lambda g_1(x_1, x_2, \dots, x_r) + \\
 & + \mu g_2(y_1, y_2, \dots, y_c).
 \end{aligned}$$

Die für das Extrem nötigen Bedingungen sind:

$$(1) \quad \frac{\partial \varphi}{\partial x_i} = \sum_j \frac{n_{ij}}{n} y_j + 2\lambda \frac{n_{.i}}{n} x_i = 0 \quad (i = 1, 2, \dots, r)$$

$$(2) \quad \frac{\partial \varphi}{\partial y_j} = \sum_i \frac{n_{ij}}{n} x_i + 2\lambda \frac{n_{.j}}{n} y_j = 0 \quad (j = 1, 2, \dots, c)$$

Aus den Beziehungen (1) und (2) folgt die Gültigkeit der Gleichungen

$$\sum_i \frac{\partial \varphi}{\partial x_i} x_i = 0, \quad \sum_j \frac{\partial \varphi}{\partial y_j} y_j = 0,$$

aus denen sich nach dem Einsetzen

$$\sum_i \sum_j \frac{n_{ij}}{n} x_i y_j + 2\lambda D(X) = 0,$$

$$\sum_i \sum_j \frac{n_{ij}}{n} x_i y_j + 2\mu D(Y) = 0$$

oder

$$\text{corr}(XY) = R = -2\lambda D(X) = -2\mu D(Y)$$

ergibt. Bei standardisierten Größen (d.h. $D(X) = D(Y) = 1$) muß dann gelten

$$(3) \quad \text{corr}(XY) = R = -2\lambda = -2\mu.$$

Aufgrund von (3) kann man die Bedingungen (1) und (2) in Form eines linearen Gleichungssystems schreiben:

$$(4) \quad \sum_j n_{ij} y_j = R n_{.i} x_i, \quad (i = 1, 2, \dots, r),$$

$$(5) \quad \sum_i n_{ij} x_i = R n_{.j} y_j, \quad (j = 1, 2, \dots, c).$$

Es kann bewiesen werden (Kendall 1973), daß das lineare Gleichungssystem (4) und (5) eine nicht triviale Lösung (ungleich Null) nur für die Werte R besitzt, die den Wurzeln der Eigenwerte der Matrix entsprechen, die durch eine einfache Transformation der ursprünglichen Kontingenztafel entsteht.

Für die Existenz einer nicht trivialen Lösung (4) und (5) muß konkret gelten:

$$(6) \quad NN'u = R^2 u,$$

wobei N die Matrix mit den Elementen

$$(7) \quad n_{ij}^* = \frac{n_{ij}}{\sqrt{n_{.i} \cdot n_{.j}}}, \quad (i = 1, 2, \dots, r), \quad (j = 1, 2, \dots, c)$$

N' die transponierte Matrix N ,
 u der Spaltenvektor mit den Komponenten

$$x_i \sqrt{n_{.i}} \quad (i = 1, 2, \dots, r) \text{ ist.}$$

Aus der Beziehung (6) ist ersichtlich, daß für die Existenz einer nicht trivialen Lösung die Gültigkeit der Gleichung

$$\det[NN' - R^2 J] = 0 \quad (J: \text{Einheitsmatrix})$$

nötig ist, d.h. R^2 ist der Eigenwert der Matrix NN' .

Der Rang der Matrix N beträgt maximal

$$m = \min(r, c),$$

so daß der Rang der Matrix NN' ebenfalls höchstens m ist. Im allgemeinen hat die Matrix NN' m verschiedene Eigenwerte

$$\lambda_0, \lambda_1, \dots, \lambda_{m-1}, \text{ wobei immer } \lambda_0 = R^2 = 1$$

einen Eigenwert der Matrix NN' darstellt (und zwar den größten), weil das Gleichungssystem (4) und (5) für $R = 1$ immer eine Lösung hat, die ungleich Null ist:

$$x_i = K, y_j = K \quad (i = 1, 2, \dots, r), \quad (j = 1, 2, \dots, c);$$

mit Rücksicht auf die Bedingungen $D(X) = D(Y) = 1$ bekämen wir aber

$$x_i = 1 \text{ und } y_j = 1.$$

Diese Werte hängen von der ursprünglichen Kontingenztafel überhaupt nicht ab, ihre Kenntnis beinhaltet demnach keine Information. Der zu dem Einheitseigenwert zugehörige Eigenvektor u^0 hat die Komponenten $\sqrt{n_{i.}}$, bzw. der normierte Eigenvektor hat die Komponenten $\sqrt{n_{i.}}/\sqrt{n}$, die nur über die relativen Größen der Marginalwerte Auskunft geben.

Die weiteren Eigenwerte, in nicht steigender Sequenz angeordnet:

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{m-1}$$

ergeben zusammen mit den ihnen entsprechenden Eigenvektoren $u^1, u^2, \dots, u^{m-1}$ schon eine nützliche Information.

Da die Eigenvektoren die Komponenten

$$u_i = x_i \sqrt{n_{i.}}$$

haben, bekommen wir aufgrund von u^1

$$(8) \quad x_i = u_i^1 / \sqrt{n_{i.}}.$$

Die Werte y_j bekommen wir dann unter Ausnutzung der Beziehung (5), oder als

$$(9) \quad y_j = \frac{1}{R n_{.j}} \sum_i n_{ij} x_i \quad (j = 1, 2, \dots, c).$$

In der gleichen Weise kann man zeigen, daß auch die Beziehung

$$N' N v = R^2 v$$

gilt, wo v ein Vektor mit den Komponenten $y_j \sqrt{n_{.j}}$ ist. Rechnerisch ist es günstig, die Aufgabe des Findens der Eigenwerte und Eigenvektoren für diejenige der Matrizen NN' und $N'N$ zu lösen, die die kleinere Ordnung hat, und die Werte der übrigen kanonischen Größen mit Hilfe von (5) oder (6) zu ermitteln.

Nun zeigen wir einige Verbindungen der beschriebenen Technik zu "traditionelleren" Verfahren der Analyse von Kontingenztafeln. Die Diagonalelemente der Matrix NN' sind gleich

$$NN'_{ii} = \sum_j \frac{n_{ij}^2}{n_{i.} n_{.j}},$$

also gilt für die Spur der Matrix NN'

$$\text{tr} NN' = \sum_i \sum_j \frac{n_{ij}^2}{n_{i.} n_{.j}}.$$

Da

$$\chi^2 = \sum_i \sum_j \frac{(n_{ij} - n_{i.} n_{.j} / n)^2}{n_{i.} n_{.j} / n} = n \left\{ \sum_i \sum_j \frac{n_{ij}^2}{n_{i.} n_{.j}} - 1 \right\}$$

ist, und die Spur der Matrix die Summe der Eigenwerte darstellt,

können wir die Größe χ^2 in Form von

$$(10) \quad \chi^2 = n(\lambda_1 + \lambda + \dots + \lambda_{m+1})$$

ausdrücken.

Die Signifikanz der Größe χ^2 ermitteln wir mit Hilfe der üblichen Tabellen der Chi-Quadrat Verteilung für ein im voraus gewähltes Signifikanzniveau und für die Anzahl der Freiheitsgrade

$$FG = (c - 1)(r - 1).$$

DIE ALGORITHMISIERUNG VON BERECHNUNGEN UND DIE GRUNDEIGENSCHAFTEN DER KANONISCHEN VARIABLEN

Die Berechnung der Werte x_i und y_j , die den einzelnen Kategorien der Kontingenztafel zugeordnet werden, benötigt einen solchen Umfang an Rechenarbeiten, daß es unerlässlich ist, auch für Tafeln kleineren Ausmaßes die EDV zu nutzen. Die Berechnung der Eigenwerte und Eigenvektoren der symmetrischen Matrix NN' , bzw. $N'N$, ist am anspruchsvollsten. Zu diesem Zweck kann man zum Beispiel einen der von Wilkinson angegebenen Algorithmen empfehlen (Wilkinson 1976). Im erarbeiteten Programm² wurde die QR-Methode verwendet, die sich gut bewährt hat.

Wir zeigen einzelne Rechenschritte und einige Eigenschaften der Werte x_i und y_j anhand eines einfachen Beispiels, das auf verschiedene Weisen modifiziert wird.

Die Kontingenztafel (Tab. 2) enthält die Ausgangsdaten für die kanonische Analyse. Die Matrix N (sie entsteht aus Tab. 1 mit Anwendung der Beziehung (7)) hat die Form

$$N = \begin{bmatrix} 0.6227524 & 0.0668153 & 0.2425356 \\ 0.3636670 & 0.5518254 & 0.1669245 \\ 0.0862796 & 0.2715377 & 0.6160411 \end{bmatrix}$$

und die Matrix NN' beinhaltet die Werte

$$N = \begin{bmatrix} 0.4511083 & 0.3038300 & 0.2212856 \\ 0.3038300 & 0.4646288 & 0.2840508 \\ 0.2212856 & 0.2840508 & 0.4606835 \end{bmatrix}$$

Tabelle 2.

	1	2	3	Σ
1	110	10	40	160
2	70	90	30	190
3	15	40	100	155
Σ	195	140	170	505

Die Eigenwerte der Matrix NN' sind:

$$\lambda_0 = 1.0000, \lambda_1 = 0.2366, \lambda_2 = 0.1398.$$

Eigenvektoren sind:

u^0	u^1	u^2
0.5629	0.6328	0.5318
0.6134	0.1114	-0.7819
0.5540	-0.7663	0.3254

Aus dem Vektor u^1 ermitteln wir durch Ausnutzung der Beziehung (8) die Werte der Variablen x_i ($i = 1, 2, 3$):

$$x_1 = 0.05, \quad x_2 = 0.0081, \quad x_3 = -0.0615.$$

Die Werte der Komponenten von Eigenvektoren sind, bis auf eine multiplikative Konstante, immer festgelegt. Die errechneten Werte sind üblicherweise normiert, d.h. zum Beispiel

$$\sum_{i=1}^3 (u_i^0)^2 = 0.5629^2 + 0.6134^2 + 0.5540^2 = 1.$$

Für die Gültigkeit der Beziehungen $D(X) = D(Y) = 1$, d.h. speziell für

$$D(X) = \sum_{i=1}^r \frac{n_{i.}}{n} x_i^2 = 1$$

ist es notwendig, daß

$$\sum_{i=1}^r \frac{n_{i.}}{n} \frac{(u_i^1)^2}{n_{i.}} = \sum_{i=1}^r \frac{(u_i^1)^2}{n} = 1$$

gilt, so daß die Komponenten der normierten Eigenvektoren mit dem Faktor $\sqrt{n}$ multipliziert werden müssen, bzw. daß die kanonischen Variablen, die aus normierten Vektoren erhalten worden sind, mit dieser Konstante multipliziert werden müssen. Nach dieser Berichtigung bekommen wir die Werte der kanonischen Variablen

$$x_1 = 1,1236, \quad x_2 = 0,1820, \quad x_3 = -1,3820.$$

Die Werte der Variablen y_j ermitteln wir aus der Beziehung (5), konkret:

$$y_1 = 1,2180, \quad y_2 = -0,4067, \quad y_3 = -1,0629.$$

Den Wert χ^2 bestimmen wir entsprechend der Beziehung (10)

$$\chi^2 = n(\lambda_1 + \lambda_2 + \dots + \lambda_{m-1}) = 505 (0,2366 + 0,1398) = 190,1$$

was mit dem auf "klassische" Weise berechneten Resultat

$$\chi^2 = \sum_i \sum_j \frac{(n_{ij} - o_{ij})^2}{o_{ij}}$$

übereinstimmt.

Der Korrelationskoeffizient ist bei den Werten x_i und y_1 für die Angaben in der Tabelle 2 gleich $R^2 = 0,2366 = 0,4864^2$, was wiederum mit dem durch eine direkte Berechnung erhaltenen Wert,

in dem jedes Wertpaar (x_i, y_j) mit der Häufigkeit n_{ij} vertreten ist, übereinstimmt.

Die Beziehungen (4) und (5), die von den Werten x_i und y_j erfüllt werden, können bei ihrer Interpretation helfen. Aufgrund der Gleichungen (4) und (5) gilt nämlich

$$\sum_j \frac{n_{ij}}{n_{i.}} y_j = R x_i \quad (i = 1, 2, \dots, r)$$

$$\sum_i \frac{n_{ij}}{n_{.j}} x_i = R y_j \quad (j = 1, 2, \dots, c),$$

wobei $\frac{n_{ij}}{n_{i.}}$ und $\frac{n_{ij}}{n_{.j}}$ Elemente der Matrizen der relativen Häufigkeiten sind (zeilen- bzw. spaltenweise relativiert).

Für die Kontingenztafel (tab. 2) ist die Matrix der relativen "Zeilenhäufigkeiten" Tab. 3 und die Matrix der relativen "Spaltenhäufigkeiten" Tab. 4.

Tabelle 3.

	11	12	13	Σ
1	0,6875	0,0625	0,2500	1,000
2	0,3684	0,4737	0,1579	1,000
3	0,0968	0,2581	0,6451	1,000

Tabelle 4.

	1	2	3
1	0,5641	0,0714	0,2353
2	0,3590	0,6429	0,1765
3	0,0769	0,2857	0,5882
Σ	1,0000	1,0000	1,0000

Die Beziehungen (4) und (5) können wir dann symbolisch so ausdrücken:

$$y_1 \begin{bmatrix} \text{1. Spalte} \\ \text{Tab. 3} \end{bmatrix} + y_2 \begin{bmatrix} \text{2. Spalte} \\ \text{Tab. 3} \end{bmatrix} + y_3 \begin{bmatrix} \text{3. Spalte} \\ \text{Tab. 3} \end{bmatrix} = R \cdot \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$$

$$x_1 [1. \text{ Zeile Tab. 4}] + [2. \text{ Zeile Tab. 4}] + [3. \text{ Zeile Tab. 4}] = R \cdot [y_1, y_2, y_3],$$

wobei

$$R = \{x_i\}^{\max}, \{y_j\} \cdot \sum_i \sum_j n_{ij} x_i y_j$$

ist.

Aus bereits oben angeführten Überlegungen ist ersichtlich, daß die Werte der kanonischen Variablen x_i und y_j vom linearen Wachstum der Häufigkeiten in der Kontingenztafel unabhängig sind, denn sie sind lediglich von relativen Häufigkeiten abhängig. Führen wir eine Permutation in der Reihenfolge der Zeilen und Spalten der Tabelle durch, so werden übereinstimmend die Reihenfolgen der Werte x_i und y_j verändert.

Vertauschen wir, zum Beispiel bei der Kontingenztafel mit den Werten der Tabelle 2, die zweite Zeile mit der ersten und gleichzeitig die zweite Spalte mit der ersten, was den Veränderungen in der Reihenfolge der Merkmalskategorien entsprechen würde, so bekommen wir die Kontingenztafel (Tab. 5).

Tabelle 5.

	1	2	3	Σ
1	90	70	30	190
2	10	110	40	160
3	40	15	100	155
Σ	140	195	170	505

Für die Angaben in Tab. 5 bekommen wir folgende Werte:

$$\lambda_0 = 1,0000, \quad \lambda_1 = 0,2366, \quad \lambda_2 = 0,1398$$

$$x_1 = 0,1820 \quad y_1 = -0,4067$$

$$x_2 = 1,1236 \quad y_2 = 1,2180$$

$$x_3 = -1,3820 \quad y_3 = -1,0629$$

Aus einem Vergleich dieser Werte mit den Ergebnissen für Tab. 2 ist ersichtlich, daß sich die Eigenwerte nicht verändert haben und daß bei den Werten x_i und y_j nur eine Veränderung der Reihenfolge zustande gekommen ist:

$$x_1 \longleftrightarrow x_2, \quad y_1 \longleftrightarrow y_2.$$

Die χ^2 -Werte ebenso wie die maximalen Korrelationskoeffizienten für die Tabellen 2 und 5 stimmen überein.

Durch eine Analyse dieser Ergebnisse gelangen wir beispielsweise zu dem Schluß, daß die Merkmale in der Tab. 2 (vor der Permutation) im Sinne der linearen Regression enger miteinander verknüpft sind, auch wenn der χ^2 -Wert für beide Tabellen gleich ist.

Konkrete Beispiele für die Anwendung der kanonischen Analyse von Kontingenztafeln findet der Interessent im Artikel von Charvát, Linhart und Vecerník (1978).

Anmerkungen:

- 1 Im Aufsatz wird die im Artikel Charvát, Linhart, Večerník (1978) verwendete Methode beschrieben.
- 2 Das Programm wurde für den Rechner NE 4120 zusammengestellt. Dieser Rechner wurde im Rechenzentrum von VSE in Prag installiert.

Literatur:

- CHARVÁT, F., LINHART, J., VEČERNÍK, J., Sociální mobilita v socialistických a kapitalistických zemích. Sociologický časopis 1978, č. 1, 40
- KENDALL, M.G., STUART, A., The Advanced Theory of Statistics. Vol. 2: Inference and Relationship. London, Griffin, 1973
- LANCASTER, H.O., Some Properties of the Bivariate Normal Distribution Considered in the Form of a Contingency Tables. Biometrika 44, 1957
- WILKINSON, J.H., REINSCH, C., Handbook for Automation Computation. Linear Algebra. Berlin, Springer Verlag 1973

CURRENT BIBLIOGRAPHY

ABBREVIATIONS

- CML Zampolli, A., Calzolari, N. (Eds.): Computational and Mathematical Linguistics. Proceedings of the International Conference on Computational Linguistics, Pisa, 27.8. to 1.9.1973, Firenze: Olschki, 1980.
- NTI Naučno-texničeskaja informacija. Serija 2: Informacionnye processy i sistemy
- PSML Prague Studies in Mathematical Linguistics 7, 1981.
- Symposium Symposium "Computational Linguistics and Related Topics", Tallinn, November 24-26, 1980. Summaries. Tallinn 1980.
- Trudy 5,6,7 Trudy po lingvostatistike 5, 6 (1980); 7 (1981), Tartuskij gosudarstvennyj universitet, Tartu.

GENERAL

1. ARAPOV, M.V.: Teorija sistem i izučenie estestvennogo jazyka: ponjatija simmetrii i statusa [Systems theory and natural language investigation: concepts of symmetry and status]. Gvišiani, D.M. (Ed.): Sistemnye issledovaniya. Metodologičeskie problemy, 1981. Moskva: Nauka 1981, 121-141.
2. ARAPOV, M.V., KRYLOV, Y.K.: Mathematical models of classification in application to some problems of statistical linguistics. Symposium, 14-16.
3. COLLIANDER, P.: Methoden zur Berechnung des Informationswerts einer in der Linguistik einsetzbaren Matrix. Kopenhagener Beiträge zur Germanistischen Linguistik 16, 1980, 1-33.
4. DZHUBANOV, A., KHASANOV, B.: Computational description of the Kazakh language. CML, Vol. 1, 75-77.

5. GERŠIĆ, S.L., ALTMANN, G.: Laut - Silbe - Wort und das MEN-
ZERATHsche Gesetz. Wodarz, H.W. (Ed.): Frankfurter pho-
netische Beiträge III, Hamburg 1980, 115-123.
6. GOEBL, H.: Typologia quantitativa oder Così fan tutte.
Grazer Linguistische Studien 11/12, 1980, 103-117.
7. KLAIVINA, S.: Statistika valodniecībā: Mācību līdžeklis
[Statistics in Linguistics]. Riga 1980, 145 pp.
8. LECOMTSEV, Y.K.: On some rules for linguistic polarity.
Symposium, 59-61.
9. MOSKOVICH, W., CAPLAN, R.: Distributive-statistical tech-
niques in linguistic and literary research. Advances
in Computer-Aided Literary and Linguistic Research,
Birmingham 1979, 245-263.
10. MULLER, C.: Langue française et linguistique quantitative.
Genève: SLATKINE Reprints, 1979, 470 pp.
11. NOSENKO, I.A.: Načala statistiki dlja lingvistov [The
Principles of Statistics for Linguists]. Moskva: Vys-
šaja škola, 1981, 157 pp.
12. PAPP, F.: Lingvostatistika i vengerskij jazyk [Linguistic
statistics and the Hungarian language]. Trudy 5, 15-37.
13. POPOVA, G.V.: On the practical use of reliability criteria
in manual and automatic abstracting. Symposium, 91-93.
14. SAUKKONEN, P.: Statistical linguistic research in Finland.
Trudy 5, 5-14.
15. ŠUBIK, S.A.: Statističeskie metody v lingvistike [Stati-
stical methods in linguistics]. Piotrovskij, R.G. et al.
(Eds.): Statistika reči i avtomatičeskij analiz teksta,
Leningrad: Nauka, 1980, 52-63.
16. SUKHOTIN, B.V.: Experiments on recognizing linguistic ob-
jects by means of algorithms of linguistic deciphering.
Symposium, 102-103.
17. TĚŠITELOVÁ, M.: Využití statistických metod v gramatice
[The Use of Statistical Methods in Grammar]. Praha:
Academia, 1981, 220 pp.

18. TUZOV, V.A.: Matematičeskaja model' jazyka [A Mathematical
Model of Language]. Leningrad 1980, 46 pp.
19. WILDGEN, W.: Differentielle Linguistik. Entwurf eines Mo-
dells zur Beschreibung und Messung semantischer und
pragmatischer Variation. Tübingen: Niemeyer 1977, 293
pp.

PHONOLOGY

20. ČIKADZE, V.G.: Statističeskij analiz fonem v odnosložnyx
slovax tureckogo jazyka [The statistical analysis of
phonemes in one-syllable words in Turkish]. Soobščenijsa
AN Gruzinskoj SSR 96, 1979, 201-209.
21. DŽUBANOV, A.K.: K voprosu o grafemnoj statistike kazaxsko-
go teksta [On the problem of graphemic statistics in
a Kazakh text]. Voprosy kazaxskoj fonetiki i fonologii,
Alma-Ata 1979, 79-86.
22. FEER, B.B.: Distribucija soglasnyx v tonal'nom ketskom ja-
zyke (Pakulixinskij govor) [The distribution of conso-
nants in tonal Ket language (Pakulixinsk dialect)]. Fo-
netika sibirskix jazykov, Novosibirsk 1970, 98-120.
23. GLUŠKOVA, I.P.: Fonologičeskaja paradigmatica i sintagmati-
ka jazyka maratxi [The Phonological Paradigmatic and
Syntagmatic of Marathi]. Avtoreferat kandidatskoj dis-
sertacii, Moskva 1980, 26 pp.
24. GRIGORIEV, V.I.: Frequency distribution of letters and
their ranks in a running text. Symposium, 43-47.
25. HORECKÝ, J., NEMCOVÁ, E.: The use of entropy in evaluating
the degree of completeness in the phonological calcu-
lus. PSML, 47-58.
26. HUG, M.: La distribution des phonèmes en français. Die
Phonemverteilung im Deutschen. Essais statistiques.
Genève: SLATKINE Reprints, 1979, 192 pp.

27. ISENGEL'DINA, A.A.: Tipologija sočetańij soglasnyx kazaxskogo i russkogo jazykov v ix forme i funkcii [A typology of the alternations of consonants in the Kazakh and Russian languages in their form and function]. Voprosy kazaxskoj fonetiki i fonologii, Alma-Ata 1979, 99-130.
28. KAASIK, Ü., TULDAVA, Ju.: Sõnalõpu grafeemotaktikast eesti-keelses tekstis. [On the graphemotactics of word endings in Estonian texts]. Trudy 5, 51-98.
29. KNOWLES, F.: The quantitative syntagmatic analysis of the Russian and Polish phonological systems. CML, Vol. 1, 79-81.
30. KURENOV, S., GOKLENOV, Dž.: Količestvennaja xarakteristika bezudarnyx glasnyx v mnogosložnyx slovax turkmenskogo jazyka [Quantitative characteristics of unstressed vowels in polysyllabic words in Turkmen]. Izvestija An Turkmenskoj SSR. Serija obščestvennyx nauk 1979/6, 50-56.
31. LIEPA, E.: Vokālisma un zilbju kvantitāte latviešu literārajā valodā [The Quantity of Vocalism and Syllables in the Latvian Literary Language]. Riga 1979.
32. PETRJANIKINA, V.I.: Častotnyj komponent intonacii kommunikativnyx tipov vyskazyvanij v jazyke kikongo [The frequency component of the intonation of communicative types of utterances in Kikongo language]. Intonacija i fonologija, Moskva 1980, 56-66.
33. PETRJANIKINA, V.I.: Rol' častotnogo komponenta intonacii v differenciacii kommunikativnyx tipov vyskazyvanij v russkom jazyke i jazykax Afriki [The role of the frequency component of intonation in the differentiation of communicative types of utterances in Russian and in African languages]. Intonacija i fonologija. Moskva 1980, 3-12.
34. ROCZAWSKI, B.: System fonostatystyczny współczesnego języka polskiego [The Phonostatistical System of Contemporary Polish]. Wrocław: Zakład Narodowy im. Ossolińskich. Wydawnictwo Polskiej Akademii Nauk, 1981, 178 pp.

35. SABOL, J.: The prosodic structure of the Slovak word. PSML, 59-80.
36. SENDLMEIER, W.F.: Der Einfluß der Quantität und Qualität auf die Perzeption betonter Vokale des Deutschen. Phonetica 38, 1981/5-6, 291-308.
37. TULDAVA, Ju.: Eesti keele sõnavara foneetilis-grafeemilised mõõted [Phonetic-Graphemic measures of Estonian lexics]. Trudy 5, 51-98.
38. VEENKER, W.: Zur phonologischen Statistik der čeremisschen (marischen) Schriftsprachen. Sovremennoe finno-ugrovedenie, Tallin 1980/2, 106-134.
39. VERNER, G.N., TAMBOVCEV, Ju.A.: Nekotorye rezul'taty fonostatističeskogo analiza zvukovogo sostava juguskogo jazyka [Some results of the phonostatistical analysis of the tonal inventory of the Yug language]. Teoretičeskie voprosy fonetiki i grammatiki jazykov narodov SSSR, Novosibirsk 1979, vypusk 1, (Ėksperimental'nye issledovanija), 47-52.

M O R P H O L O G Y

40. ARAPOV, M.V., HERZ, M.M.: Frequency and age as characteristics of a word. CML Vol. 1, 3-7.
41. BERGER, T.: The formation of the imperative in modern Russian. Russian Linguistics 6, 1981, 65-80.
42. GRYAZNUKHINA, T.A., KOMAROVA, L.I., KRITSKAYA, V.I.: Functional peculiarities of the noun in scientific abstracts. Symposium, 48-51.
43. HOOGE, D.J.: O nekotoryx kvantitativnyx aspektax upotreblenija modal'nyx glagolov (na materiale nemeckogo jazyka) [On some quantitative aspects of the use of modal verbs

in German]. *Linguistica* 13, 1980, 55-63.

44. JAKUBOVKA, N.K.: Rasstanovka affiksov v uzbekskix slovoformax [The arrangement of affixes in Uzbek word forms]. Tāskent: Izdatel'stvo "Fan" Uzbekskoj SSR, 1979, 108 pp.
45. KAASIK, Ü, TULDAVA, Ju.: Some problems of the automatic morphological analysis of word forms in Estonian texts. Symposium, 54.
46. PETRAVIČIUS, L.: Zur Valenz des Verbs in der deutschen wissenschaftlichen Periodik des 19. - 20. Jh. *Kalbotyra* 31, 1980/4, 35-53.
47. POLJUŽKIN, M.M.: Rol' kvantitativnyx metodov v issledovanii funkcionirovaniya affiksov [The role of quantitative methods in the study of the functioning of affixes]. Slovoobrazovanie i ego mesto v kurse obučeniya inostrannogo jazyka 7, Vladivostok 1979, 72-89.
48. VAJŠNORAS, A.: Produktivnost' i častotnost' funkcionirovaniya leksiko-sintaksičeskoj eljativnoj formy v sovremennyx ital'janskom i francuzskom jazykax [The productivity and frequency of the functioning of the lexical-syntactic relative form in contemporary Italian and French]. *Kalbotyra* 31, 1980/5, 7-18.

S Y N T A X

49. BEEBE, R.D.: The frequency distribution of English syntagms. *CML*, Vol. 1, 9-22.
50. NARO, A.J.: The social and structural dimensions of a syntactic change. *Language* 57, 1981, 63-98.
51. ŠTĚPÁN, J.: On the statistics of the complex sentence (The degree of dependence of subordinate clauses and thinking). *PSML*, 113-122.

S E M A N T I C S

52. ARAPOV, M.V.: Reguljarnost' semantičeskix otnošenij v estestvennom jazyke [The regularity of semantic relations in natural languages]. *NTI* 1980/9, 1-8.
53. HÉRAULT, D.: Compréhension automatique et spectre sémantique. Paris: Jean-Favard 1981, 289 pp.
54. KARAULOV, Ju.N.: Častotnyj slovar' semantičeskix množitelej russkogo jazyka [A Frequency Dictionary of Semantic Factors in Russian]. Moskva: Nauka 1980, 207 pp.
55. LECOMTSEVA, M.I., LECOMTSEV, Y.K.: Towards algebraic modeling of linguistic semantics. Symposium, 62-71.
56. MOSKOWICH, W.: Polysemy in natural and artificial (planned) languages. *Statistical Methods in Linguistics* 1, 1977, 5-28.
57. RIEGER, B.: Linguistic semantics and the problem of vagueness: on analysing and representing word meaning. *Advances in Computer-Aided Literary and Linguistic Research*, Birmingham 1979, 271-288.
58. STERNIN, I.A.: Iz nabljudenij nad verojatnostnymi semami v značenii slova [Observations on probability senses in word meaning]. *Sopostavitel'no-semantičeskie issledovaniya russkogo jazyka*, Voronež 1980, 14-20.
59. TĚŠITELOVA, M.: Sémantika a statistika [Semantics and statistics]. *Slovo a slovesnost* 51, 1980, 100-105.

L E X I C O L O G Y

60. ALEKSEEV, P.M., KAŠIRINA, M.E., TARASOVA, E.M. (Eds.): Častotnyj anglo-russkij fizičeskij slovar'-minimum [A Frequency English-Russian Minimum Physics Vocabulary]. Moskva: Voenizdat, 1980, 288 pp.

61. ALEKSEEV, P.M., KOPYLOVA, M.E.: O serijnyx dvujazyčnyx učebnyx častotnyx slovarjax [On serial bilingual frequency dictionaries for school use]. Trudy 6, 3-13.
62. ARZIKULOV, X.A., ARZIKULOVA, Ž.X.: O dvujazyčnyx častotnyx slovarjax-minimumax [On bilingual frequency vocabularies]. Naučnye doklady vyššej školy Filologičeskie nauki, Moskva 1980/3, 78-82.
63. AZAROVA, I.V., GOROKHOVA, S.I., KUZNETSOVA, E.L.: Machine-produced word indexes and concordances. Symposium, 17-18.
64. BRUNET, E.: Le vocabulaire de Jean Giraudoux. Structures et évolution. Genève: SLATKINE Reprints, 1978, 681 pp.
65. DOLPHIN, B.: Vocabulaire et lexique. Modèles mathématiques pour une linguistique quantitative. Genève: SLATKINE Reprints, 1980, 156 pp.
66. DUGAST, D.: Vocabulaire et discours. Essai de lexicométrie organisationnelle. Fragments de lexicologie quantitative. Avec trois séries de 730 couples de données rendant compte en 18 pages du roman de Maupassant "Fort comme la mort". Genève: SLATKINE Reprints, 1979, 137 pp.
67. EGIAZARJAN, Ė.V.: Nekotorye voprosy postroenija avtomatičeskix slovarj dlja mašinno perevoda [Some problems in the construction of automatic vocabularies for machine translation]. Strukturnyj analiz teksta. Erevan 1979, 79-89.
68. GERD, A.S.: A computer-aided terminological data bank in the light of the theory of scientific and technical lexicography. Symposium, 29-30.
69. GRIGORIEV, H.G., MARUSENKO, M.A.: An algorithm of compiling a terminological dictionary word-list by a computer. Symposium, 39-41.

70. GRIGOREVA, A.S.: O častotnom slovare russkoj obixodnoj pis'mennoj reči [On a frequency dictionary of everyday written Russian]. Trudy 6, 24-31.
71. GUROVA, N.V.: Computational parameters of basic vocabulary selection from a word count. Symposium, 52-53.
72. JOHNSON, R.L.: Measures of vocabulary diversity. Advances in Computer-Aided Literary and Linguistic Research, Birmingham 1979, 213-217.
73. KVJATNOVSKAJA, T.A.: O dvujazyčnyx častotnyx slovarjax [On bilingual frequency dictionaries]. Izvestija Akademii Nauk Latvijskoj SSR, Riga 1981/2, 66-70.
74. LECOMTSEV, Y.K.: Notes on an algebraic model of lexical fields. Symposium, 55-58.
75. MALAXORSKIJ, L.V.: Principy častotnoj stratifikacii slovarnogo sostava jazyka [The Principles of Frequential Stratification of a Language's Vocabulary]. Leningrad: Nauka 1980, 223 pp.
76. MAMEDOVA, M.G., SKOROKHODKO, E.F.: Computerized network analysis of lexical systems. Symposium, 72-73.
77. MELIKJAN, N.A.: Častotnyj mikroslovar' pod-jazyka telegramm [A frequency micro-dictionary of the sublanguage of telegrams]. Strukturnyj analiz teksta. Erevan 1979, 97-116.
78. MULLER, C.: Le vocabulaire du théâtre de Pierre Corneille. Étude de statistique lexicale. Genève: SLATKINE Reprints, 1979, 379 pp.
79. MURAKI, S.: Comparison of basic vocabularies of Japanese and German [In Japanese]. Mathematical Linguistics 12, 1981, 356-366.
80. NAKANO, H.: The number of words in the word list by semantic principles [In Japanese]. Mathematical Linguistics 12, 1981, 376-381.

81. SEEGBRECHT, A.-L.: Wortfrequenzvergleiche zwischen dem Čechischen und Finnischen. Dissertation. Freiburg i. Br. 1981.
82. TANAKA, Y.: Analysis of technical terms [In Japanese]. Mathematical Linguistics 12, 1981, 367-376.
83. TULDAVA, J.: A method for calculating vocabulary growth as a function of text length. Symposium, 110-114.
84. TULDAVA, Ju.: O teoretiko-metodologičeskix osnovax kvantitativno-sistemnogo analiza leksiki (I) [On the theoretical-methodological foundations of the quantitative-systemic lexicon analysis]. Linguistica 13, 1980, 143-158.
85. TULDAVA, Ju.: O teoretiko-metodologičeskix osnovax kvantitativno-sistemnogo analiza leksiki (2): Lingvističeskie aspekty issledovanija [On the theoretical methodological foundations of the quantitative-systematic lexicon analysis (2): The linguistic aspects of the research]. Linguistica 14, 1981, 114-133.
86. TUŠIŠVILI, M.A.: Nekotorye metody statističeskogo analiza vremennoj struktury rečevogo signala [Some statistical methods for analyzing the tense structure of speech signals]. Trudy Instituta sistem upravljenija Akademii Nauk GSSR, Tbilisi, 19, 1980/2, 144-153.
87. VIKS, Ü.: An automatic system for generation of dictionaries. Symposium, 119-120.
88. XOFFMANN, L.: Častotnye slovari pod "jazykov nauki i tekhniki - ključ k ponimaniju special'nyx tekstov [Frequency dictionaries of scientific and technological sublanguages - a key to the understanding of technical texts]. Trudy 6, 145-152.

TEXT ANALYSIS

89. ALEKSEEV, P.M.: Computational linguistics and quantitative typology of text. Symposium, 5-7.
90. ALEKSEEV, P.M.: O kvantitativnoj tipologii teksta [On quantitative typology of text]. Trudy 7, 3-13.
91. ALEXandrova, L.N.: The postulates of Propp-Revzin and the repeat rate of parts of speech in texts samples. Symposium, 8-10.
92. ARAPOV, M.V.: Sistemnyj analiz leksičeskoj struktury tekstov [The systems analysis of the lexical structure of texts]. Gvišiani, D.M. (Ed.): Sistemnye issledovanija. Metodologičeskie problemy, 1980. Moskva: Nauka 1981, 372-402.
93. BAILEY, R.W.: Authorship attribution in a forensic setting. Advances in Computer-Aided Literary and Linguistic Research, Birmingham 1979, 1-20.
94. BEKTAEV, K.B.: Ėntropija tjurkskix pečatnyx tekstov [The entropy of Turkic printed texts]. Voprosy kazaxskoj fonetiki i fonologii, Alma Ata 1979, 43-52.
95. BELETSKAJA, I.P.: A knowledge storage system in dialogue mode in natural language. Symposium, 19-21.
96. BORODA, M.G.: Ob odnoj količestvennoj zakonomernosti ritmiki vokal'noj melodii [On a quantitative regularity in the rhythm of vocal melody]. G. Toradze (Ed.): Konferencija molodyx jazykovedov. Tezisy i materialy. Tbilisi 1980, 18-21.
97. BORODA, M.G., NADAREISHVILI, G.Sh.: On certain statistical characteristics of relations between text and melody in vocal music. Symposium 22-26.
98. BUSA, R.: How quantitative information on words, forms, and lemmas is presented in the Index Thomisticus. CML, Vol. II 887-892
99. DOLPHIN, C.: Méthodes de la statistique linguistique et vocabulaire fantastique de Malpertuis. Genève: SLATKINE

Reprints, 1980, 301 pp.

100. DUGAST, D.: Vocabulaire et stylistique. Théâtre et dialogue. Genève: SLATKINE Reprints, 1980, 293 pp.
101. FREY, E.: Subjective word frequency estimates and their stylistic relevance in literature. *Poetics* 10, 1981, 395-407.
102. GASPAROV, M.L.: Ital'janskij stix: sillabika ili sillabotnika? (Opyt izpol'zovanija verojatnostnyx modelej v sravnitel'nom stixovedenii) [Italian verse: syllabic or syllabo-tonic? (An attempt to use probability models in comparative verse research)]. Grigor'ev, V.P. (Ed.): Problemy strukturnoj lingvistiki 1978. Moskva 1981, 198-218.
103. GINDIN, S.I.: Ritmika, intonacija i smyslovaja kompozicija v poëme VI. Lugovskogo "Kak čelovek plyl s Odisseem" [Rhythm, intonation, and semantic composition in VI. Lugovskij's epic poem "Kak čelovek plyl s Odisseem"]]. Grigor'ev, V.P. (Ed.): Problemy strukturnoj lingvistiki 1978. Moskva 1981, 230-265.
104. GRIDNEVA, L.M., DARCHUK, N.P., MURAVITSKAYA, M.P., PEREBYNOS, V.I.: Symmetry elements in texts of scientific abstracts. Symposium, 35-38.
105. GROŠEVA, A.V.: Statističeskaja odnorodnost' tekstov na sintaksičeskom urovne [The statistical homogeneity of texts on the syntactic level]. Piotrovskij, R.G. et al. (Eds.): Statistika reči i avtomatičeskij analiz teksta, Leningrad: Nauka, 1980, 72-98.
106. IBRAGIMOV, S.I.: O nekotoryx strukturno-verojatnostnyx parametrah kirgizskogo naučno-texničeskogo teksta [On some structural-probabilistic parameters of a scientific-technical text in Kirghiz]. Sovetskaja tjurkologija 1979/5, 72-79.
107. JAKUBAJTIS, T.A., SKLJAREVIČ, A.N.: Korreljacionnye xarakteristiki častej reči v svjaznyx textax [Correlatio-

- nal Characteristics of Parts of Speech in Interrelated Texts]. Riga: Zinatne, 1980, 135 pp.
108. KJETSAA, G.: A quantitative norm for the use of epithets in the age of Puškin. Russian Romanticism. Studies in the Poetic Codes. Stockholm Studies in Russian Literature 10, 1979, 204-226.
109. KJETSAA, G.: Stil' i norma [Style and norm]. *Linguistica* 14, 1981, 48-67.
110. KLIMEŠ, L.: On some quantitative aspects of the sentence in Loewenbrugk's Chronicle (1685). *PSML*, 123-136.
111. KOKKOTA, V.: Častotnost' anglijskix glagol'nyx form v naučno-texničeskoj literature [Frequencies of English verb forms in scientific and technical literature]. *Trudy* 6, 59-72.
112. KRASNOVA, E.A.: Leksiko-semantičeskie osobennosti žanra (Opyt statističeskogo analiza na materiale anglijskix i šotlandskix narodnyx ballad) [Lexical and semantic distinctions of a genre (Attempt at a statistical analysis based on English and Scottish folk ballads)]. *Trudy* 6, 73-85.
113. KRAUS, J.: Ličnye mestoimenija v češskix publicističeskix tekstax (Kvantitativnyj analiz) [Personal pronouns in Czech journalistic texts (Quantitative analysis)]. *PSML*, 27-46.
114. LEVIN, Ju.I.: Zamečanija o priloženii matematičeskoj statistiki dlja izučeniya zavisimostej i svjazej meždu xarakteristikami xudožestvennyx tekstov [Notes on application of mathematical statistics to investigation of dependences and relations between parameters of literary texts]. *Trudy* 7, 46-59.
115. LIŠKOVÁ, Z.: Two signals of readability in journalistic utterance. *PSML*, 81-92.
116. MANASJAN, N.S.: O raspredelenijax terminov v anglijskom naučno-texničeskom tekste [On distribution analysis in

English scientific texts (Sublanguage of active Oscillators)]. Trudy 7, 60-73.

117. MARTEM'YANOV, Y.S.: A word-correlated syntax and text analysis. Symposium, 74.
118. MARTY'NENKO, G.Y.: Automatic classification of texts according to syntactic features. Symposium, 75-77.
119. MARUSENKO, M.A.: O izmerenii svyazi otraslevykh terminosistem s primeneniem ÈVM [Computer measuring of lexical connection of terminological systems relating to a certain branch of industry]. Trudy 7, 74-81.
120. MELETTI, F.: Proportion between number of words and number of forms in the works of the "Index Thomisticus". CML, Vol. 1, 89-98.
121. MIZUTANI, S.: Calculation of the distribution of word frequencies [In Japanese]. Mathematical Linguistics 12, 1981, 339-356.
122. NADAREIŠVILI, I.S. Statističeskij analiz morfoložičeskoj struktury teksta (Na materiale prozy K. Gamsaxurdija) [A statistical analysis of morphological text structure (Based on the prose of K. Gamsaxurdij)]. Izvestija Akademii Nauk GSSR. Serija jazykov i literatura, Tbilisi, 1980/3, 143-148.
123. NADAREIŠVILI, I.Sh., ORLOV, Y.K.: On the effect of the group of frequent words upon the frequency structure and lexical concentration of a text. Symposium, 81-86.
124. OSTAPENKO, V.E.: Paraboličeskie i giperboličeskie zakonomernosti tekstoobrazovaniya [Parabolic and hyperbolic regularities of text formation]. Trudy 6, 106-112.
125. PEREBEJNOS, V.I.: Raspredelenie glagolov v naučno-referativnom tekste [Distribution of verbs in scientific abstracts]. Trudy 7, 91-100.
126. PIOTROVSKIY, R.G. et al. (Eds.): Statistika reči i avtomatičeskij analiz teksta [The Statistics of Speech and the Automatic Analysis of Texts]. AN SSSR Institut Ja-

zykoznanija. Leningrad: Nauka, 1980, 223 pp.

127. REMMEL, M.: On text compression and rank statistics. Symposium, 94-96.
128. SAUKA, L.I.: O količestvennom analize stixopesennykh vstavok litovskikh narodnykh skazok. [On the quantitative analysis of rhymed song-inserts in Lithuanian folk tales]. Trudy AN Litovskoj SSR, Serija obščestvennykh nauk 1980/1, 117-123.
129. SLEPAK, B.A.: "Prolegomeny" k statističeskoj teorii teksta ["Prolegomena" to the statistical theory of text]. Trudy 7, 101-119.
130. SLIVNJAK, D.I.: Statističeskie xarakteristiki mežsegmentnykh granic v dialoge [Statistical characteristics of intersegmental boundaries in the dialogue]. Trudy 7, 120-135.
131. TEŠITELOVA, M.: On the language of the present-day publicist prose from the quantitative point of view. PSML, 9-26.
132. TULDAVA, Ju.A.: Opyt klassifikacii tekstov s pomošč'ju klaster-analiza [On a possibility of the realization of cluster-analysis]. Trudy 7, 136-157.
133. VALK-FALK, M.: An attempt of statistical analysis of melodic structures in two-part keyboard music. Symposium, 115-118.
134. VAŠAK, P.: Tekstologija - sistema - stemma [Textology - system - stemma]. PSML, 137-147.
135. ŽILINSKENE, V.Ju.: Statističeskij analiz morfoložii litovskikh gazetno-žurnal'nykh tekstov (po dannym častotnogo slovarja litovskikh gazetno-žurnal'nykh tekstov) [A Statistical Analysis of the Morphology of Lithuanian Magazine and Newspaper Texts, Based on a Frequency Dictionary of Lithuanian Magazine and Newspaper Texts]. Avtoreferat kandidatskoj dissertacii, Vil'njus 1979, 24 pp.

statistical approach). *Linguistica* 13, 1980, 132-142.

DIALECTOLOGY

152. KUKORIN, V.N.: Distribucija glasnyx v čalkanskom dialekte altajskogo jazyka [(The distribution of vowels in the Čalkan dialect of Altai language). *Issledovaniya zvukovyx sistem sibirskix jazykov*, Novosibirsk 1979, 14-28.
153. VERTE, L.A.: Nekotorye distributivnye xarakteristiki so-glasnyx fonem kazymskogo dialekta xantyskogo jazyka [A few distributive characteristics of consonantal phonemes of the Kazym dialect in Khanty]. *Issledovaniya zvukovyx sistem sibirskix jazykov*, Novosibirsk 1979, 144-147.

QUANTITATIVE LINGUISTICS

Appeared

- Vol. 1. Altmann, G. (Ed.), *Glottometrika 1*
 Vol. 2. Grotjahn, R., *Linguistische und statistische Methoden in Metrik und Textwissenschaft*
 Vol. 3. Grotjahn, R. (Ed.), *Glottometrika 2*
 Vol. 4. Strauß, U., *Struktur und Leistung der Vokalsysteme*
 Vol. 5. Matthäus, W. (Ed.), *Glottometrika 3*
 Vol. 6. Grotjahn, R., Hopkins, E. (Eds.), *Empirical Research on Language Teaching and Language Acquisition*
 Vol. 7. Altmann, G., Lehfeldt, W., *Einführung in die quantitative Phonologie 1*
 Vol. 8. Altmann, G., *Statistik für Linguisten 1*
 Vol. 9. Hopkins, E., Grotjahn, R. (Eds.), *Studies in Language Teaching and Language Acquisition*
 Vol. 10. Skorochod'ko, E.F., *Semantische Relationen im Lexikon und in Texten*
 Vol. 11. Grotjahn, R. (Ed.), *Hexameter Studies*
 Vol. 12. Rieger, B. (Ed.), *Empirical Semantics. A Collection of New Approaches in the Field, Vol. 1*
 Vol. 13. Rieger, B. (Ed.), *Empirical Semantics. A Collection of New Approaches in the Field, Vol. 2*
 Vol. 14. Lehfeldt, W., Strauß, U. (Eds.), *Glottometrika 4*

In Preparation

- Alekseev, P., *Statistische Lexikographie*
 Altmann, G., *Diskrete Wahrscheinlichkeitsverteilungen*
 Arapov, M.V., Cherc, M.M., *Mathematische Methoden in der historischen Linguistik*
 Boy, J., *Mathematische Grundverfahren für Linguistik*
 Boy, J., *Phonetische Dechiffrierung von Buchstaben*
 Brainerd, B. (Ed.), *Historical Linguistics*
 Frank, H. (Hrsg.), *Interlinguistik*
 Gindin, S. (Ed.), *Soviet contributions to metrics*
 Goebel, H. (Ed.), *Dialectology*
 Guiter, H., Arapov, M.V. (Eds.), *Studies on Zipf's Law*
 Lesochin, M.M., Piotrowski, R.G., *Mathematical Models in Quantitative Linguistics*
 Matthäus, W. (Ed.), *Psycholinguistics*
 Orlov, Ju.K., Boroda, M.G., Nadarejvili, I.S., *Sprache, Text, Kunst. Quantitative Analysen*
 Piotrowski, R.G., Bektaev, K.B., Piotrowskaja, A.A., *Mathematische Linguistik*

The series publishes

- Mixed volumes
- Monothematic volumes
- Textbooks
- Monographs
- Frequency dictionaries

Contributions can be sent to G. Altmann, Sprachwissenschaftliches Institut der RUB, Postfach 102148, 4630 Bochum, West Germany
Orders for individual volumes or the entire series should be directed to Studienverlag Dr. N. Brockmeyer, Querenburger Höhe 281, 4630 Bochum, West Germany.