

QUANTITATIVE LINGUISTICS

Editor G. Altmann, Bochum

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

CIP-Kurztitelaufnahme der Deutschen Bibliothek

Grotjahn, Rüdiger:

Linguistische und statistische Methoden in Metrik
und Textwissenschaft / Rüdiger Grotjahn. —
Bochum: Studienverlag Brockmeyer, 1979. —
(Quantitative Linguistics; vol.2)
ISBN 3-88339-062-3

ISBN 3-88339-062-3

Alle Rechte vorbehalten

© 1979 by Studienverlag Dr. N. Brockmeyer

Querenburger Höhe 281, 4630 Bochum 1

VORWORT

Diese Arbeit verfolgt zum einen das Ziel, die Bedeutung linguistischer und statistischer Modellbildung für Metrik und Textwissenschaft deutlich zu machen. Zum anderen sollen eine Reihe von Anwendungsmöglichkeiten der Statistik bei der Analyse formaler Textstrukturen aufgezeigt werden. Auf diese Weise soll ein Beitrag zu einer stärkeren 'Objektivierung' und 'Empirisierung' von Metrik und Textwissenschaft geleistet werden.

Die vorliegende Arbeit wurde im Herbst 1977 von der Abteilung für Philologie der Ruhr-Universität Bochum als Dissertation angenommen. Außer einigen geringfügigen Korrekturen und Kürzungen habe ich den Text unverändert gelassen.

An dieser Stelle möchte ich G. Altmann, U.-L. Figge, G. Jaeschke, M. Job, W. Lehfeldt, W. Matthäus sowie Th. Stoffer für ihre Hinweise und kritischen Stellungnahmen danken. Für verbleibende Fehler bin ich natürlich allein verantwortlich.

Bochum, im Frühjahr 1979

R.G.

I N H A L T

	Seite
VERZEICHNIS DER TABELLEN	III
VERZEICHNIS DER ABBILDUNGEN	V
1. METRIK UND LINGUISTIK	1
2. ZUR WISSENSCHAFTSTHEORETISCHEN BASIS EINER EMPIRISCHEN METRIK	5
3. EIN RAHMENMODELL ZUR ANALYSE METRISCHER TEXTE	9
4. TEXTE UND ZUFALLSFOLOGEN	17
5. RHYTHMUS UND METRUM	23
5.1. Der Rhythmus	23
5.2. Das Metrum	27
5.3. Der Vers	29
5.4. Metrum, rhythmische Struktur und rhythmische Realisation	30
6. DIE INTERDEPENDENZ VON METRISCHEN UND SPRACHLICHEN STRUKTUREN	34
6.1. Das HEMPEL-OPPENHEIM-Schema einer wissen- schaftlichen Erklärung	34
6.2. Der Kontakt zwischen metrischen und sprach- lichen Strukturen	36
6.3. Die Abhängigkeit des metrischen Systems vom Sprachsystem	37
6.4. Zur diachronen Entwicklung metrischer Systeme	39
6.5. Zur Anwendung des HEMPEL-OPPENHEIM-Schemas in der Metrik	41
7. DIE SILBE ALS METRISCHE KONSTITUENTE	44
8. DER REIM	47
8.1. Zur Definition des Reims	47
8.2. Zur Semantik des Reims	52
8.3. Reiner und unreiner Reim	54
8.4. Der unreine Reim	56
8.5. Reim und Sprachstruktur	59

9. METRIK UND SPRACHTYPOLOGIE	63
10. DIE FUNKTION METRISCHER STRUKTUREN	66
11. ZUR THEORIE DER MATHEMATISCHEN METRIK	72
11.1. Die Verwendung mathematischer Methoden in der Metrik	72
11.2. Ein Begriffsrahmen für die quantitative Textanalyse	75
11.3. Der forschungslogische Ablauf quantitativer Textanalysen	82
11.4. Skalierung und Messung textueller Eigenschaften	86
12. QUANTITATIVE ANALYSE METRISCHER TEXTE	90
12.1. Der "Erlkönig" von GOETHE als Analyseobjekt	90
12.2. Analyse der Senkungen	92
12.3. Korrelation zwischen Inhalt und Formalstruktur	114
12.4. Verstyp und Verslänge	130
12.5. Analyse der metrischen Bindung	157
12.6. Verteilung der Anzahl der Silben pro Wort	173
12.7. Analyse der Lautstruktur	179
13. DIE GENERATIVE METRIK	197
13.1. Der generative Ansatz in der Metrik	197
13.2. Generative Metrik und probabilistische Regelbewertung	201
13.3. Ein probabilistisches Modell für den lateinischen Hexameter	205
13.3.1. Regressionsanalyse	205
13.3.2. Die Ansätze BOLDRINI's und HERDAN's	210
13.3.3. Das Modell einer nicht-homogenen MARKOV-Kette	213
13.4. Probabilistische generative Metrik	224
14. ZUSAMMENFASSUNG UND AUSBLICK	227
ANMERKUNGEN	230
ANHANG	263
LITERATURVERZEICHNIS	265

VERZEICHNIS DER TABELLEN

	Seite
12.1 Häufigkeit der verschiedenen Senkungstypen vor der i-ten Hebung im "Erlkönig"	92
12.2 Beobachtete und erwartete Häufigkeiten für kurze und lange Senkungen im "Erlkönig"	95
12.3 Berechnung der Regression der langen Senkungen $\hat{p}_i$ auf den Hebungstyp x_i	99
12.4 Beobachtete und geschätzte Werte für die relativen Häufigkeiten von langen und kurzen Senkungen pro pro Vers im "Erlkönig"	99
12.5 Anteil der linearen Regression und der Abweichungen von der Regression am Gesamt- X^2 für die Daten der Tabelle 12.3	102
12.6 Häufigkeit der verschiedenen Senkungstypen vor der i-ten Hebung in "Der Totentanz" von GOETHE	102
12.7 Absolute (f.) und relative (%) Häufigkeiten von kurzen und langen Senkungen im "Erlkönig" und in "Der Totentanz"	103
12.8 Anzahl der kurzen und langen Senkungen pro Strophe im "Erlkönig"	110
12.9 Anteil der linearen Regression und der Abweichungen von der Regression am Gesamt- X^2 für die Daten der Tabelle 12.8	112
12.10 Beobachtete und geschätzte Werte für die relativen Häufigkeiten von langen und kurzen Senkungen pro Strophe im "Erlkönig"	114
12.11 Bewertung der inhaltlichen Spannung der Strophen des "Erlkönig" durch 20 Versuchspersonen	116
12.12 Rangzahlen der Zeilen-Werte aus Tabelle 12.11	117
12.13 Rangzahlen der Spalten-Werte aus Tabelle 12.11	124
12.14 Multiple Vergleiche zwischen den Rangsummen aus Tabelle 12.12	125
12.15 Korrelation zwischen der inhaltlichen Spannung und der Anzahl der langen Senkungen im "Erlkönig"	128
12.16 Häufigkeiten der verschiedenen Verstypen im "Erlkönig"	130
12.17 Anzahl der Wörter pro Vers im "Erlkönig" und in "Der Totentanz"	133
12.18 Kenngrößen der Daten aus Tabelle 12.17	136
12.19 Anzahl der Wörter pro Vers in den einzelnen Strophen des "Erlkönig"	142
12.20 Varianzanalyse der Daten aus Tabelle 12.19	143
12.21 Anzahl der Silben (S) pro Vers (V) im "Erlkönig"	144

	Seite
12.22 Zusammenfassung der Daten aus Tabelle 12.21	144
12.23 Iterationswerte der Verstypen aus Tabelle 12.21	154
12.24 Beobachtete und erwartete Anzahl der Iterationen von markierten (M) und unmarkierten Silben (U) im "Erlkönig" und in "Der Totentanz"	157
12.25 Beobachtete und erwartete Anzahl der Iterationen in der "Aeneis" und im "Bellum Gallicum"	160
12.26 Kontingenzkoeffizienten für den Grad der metrischen Bindungen im "Erlkönig"	164
12.27 Beobachtete (ϕ_b) und theoretische (ϕ_t) Werte der Kontingenzkoeffizienten des "Torquato Tasso"	171
12.28 Verteilung der Anzahl der Silben pro Wort im "Erlkönig" und in 10 Briefen GOETHEs aus dem Jahre 1782	174
12.29 Verteilung der Anzahl der Silben pro Wort in zwei lateinischen Verstexten, zwei lateinischen Prosatexten und zwei deutschen Balladen	177
12.30 Häufigkeiten der Phoneme des Deutschen nach MEIER (1967:251ff) und Häufigkeiten der Phoneme des "Erlkönig"	182
12.31 Vergleich der Phonemfrequenz im "Erlkönig" und bei MEIER (1967:253)	193
12.32 Anzahl der Konsonanten und Vokale im "Erlkönig" und bei MEIER (1967:253)	195
13.1 Absolute und prozentuale Häufigkeiten des lateinischen Hexameters entsprechend der Anzahl und der Position von Daktylen (d) und Spondeen (s) nach DROBISCH (1866)	207
13.2 Absolute und prozentuale Häufigkeiten von Daktylen und Spondeen in den ersten vier Hexameterfüßen bei VERGIL (n = 1760)	208
13.3 Schätzwerte für die Häufigkeit von Daktylen und Spondeen (%)	209
13.4 Beobachtete prozentuale Häufigkeiten ($\hat{p}_i$) und errechnete prozentuale Wahrscheinlichkeiten (p_i) der Hexametertypen nach BOLDRINI (1948) und HERDAN (1954)	211
13.5 Relative Häufigkeiten $\hat{p}_i$ und errechnete Wahrscheinlichkeiten p_i von Hexametertypen bei VERGIL	217
13.6 Anfangswahrscheinlichkeit eines Daktylus und Wahrscheinlichkeit des Übergangs von einem Daktylus zu einem Daktylus und von einem Spondeus zu einem Spondeus	220
13.7 Schätzung der Regressionsparameter für die Werte der Daktylen (y) aus Tabelle 13.6	221
13.8 Relative Häufigkeiten $\hat{p}_i$ und errechnete Wahrscheinlichkeiten p_i von Hexametertypen bei VERGIL (Regressionsmodell)	223

VERZEICHNIS DER ABBILDUNGEN

	Seite
3.1. Textwissenschaftliche Kommunikationskette	10
12.1. Notation des "Erlkönig"	91
12.2. Verlauf der Senkungen im Vers	100
12.3. Korrelogramm der metrischen Bindungen im "Erlkönig" (%)	167
12.4. GOETHE, Torquato Tasso, Jambus (u b)	168
12.5. BENN, Prosa und Szenen	168
12.6. Absolutbeträge der Korrelationskoeffizienten des "Erlkönig"	171

1. METRIK UND LINGUISTIK

1.1. Als Hauptgegenstand der Metrik¹ kann die Formalstruktur von Texten, die ich *metrische Texte*² nennen möchte, angesehen werden. Metrische Texte lassen sich in erster Annäherung dadurch charakterisieren, daß sie eine erhöhte Regularität aufweisen - und zwar erhöht im Verhältnis zu der im Rahmen der jeweiligen Sprachgesetze üblichen Regularität. Diese erhöhte Regularität kann sich z.B. äußern in einer numerisch geregelten Alternanz von betonten und unbetonten Silben oder in einer numerischen Regulierung der Anzahl der Silben innerhalb eines bestimmten syntaktischen Rahmens.

Häufig werden Texte, die formale Charakteristiken dieser Art aufweisen, auch zu den poetischen bzw. literarischen Texten gezählt. Dies ist jedoch m.E. zumindest aus zwei Gründen problematisch. Einmal weisen auch Texte, die meist nicht zu den poetischen oder literarischen Texten gerechnet werden, wie z.B. Werbetexte oder antike Gesetzestexte, zuweilen metrische Ordnungsprinzipien auf. Zum anderen ist es m.E. nicht sehr sinnvoll, die Prädikate 'literarisch' bzw. 'poetisch' ohne eine umfassende Einbeziehung der Textrezeption zu explizieren (vgl. hierzu z.B. IHWE 1972). Es scheint mir deshalb wenig angebracht, metrische Texte global als poetisch bzw. literarisch zu kennzeichnen (vgl. auch IHWE 1975:394f).³

1.2. Metrische Texte können unter einer Vielzahl von Gesichtspunkten betrachtet werden. Dies hat dazu geführt, daß metrische Texte auch von einer Reihe verschiedener Wissenschaften untersucht werden. So beschäftigen sich vor allem Literaturwissenschaft, Poetik, Linguistik und Stilistik⁴ und auch - allerdings mehr am Rande - z.B. Psychologie, Ästhetik, Semiotik, Rhetorik und Musikwissenschaft mit Fragen der Metrik.

Dies hat zur Konsequenz, daß die Metrik, je nachdem welcher Aspekt im Vordergrund der Betrachtungsweise steht, wissenschaftssystematisch verschiedenen Disziplinen zugerechnet wird. So wird die Metrik z.B. als Teildisziplin von Linguistik, Poetik⁵ (Literaturwissenschaft) oder Stilistik angesehen.

In diesem Kontext kommt m.E. jedoch der Linguistik eine besondere Bedeutung zu.⁶ Zum einen bedarf es auch bei einer literaturwissenschaftlichen oder ästhetischen Analyse eines beliebigen Textes - also auch eines metrischen Textes - einer vorgängigen Klärung seiner sprachlichen Struktur.⁷ Zum anderen hat die Linguistik z.B. in Form des Strukturalismus und der generativen Transformationsgrammatik eine z.B. im Vergleich zur Literaturwissenschaft relativ gut fundierte empirische Methodologie entwickelt (vgl. IHWE 1972:20 ff).

Es ist deshalb wenig erstaunlich, daß innerhalb solcher Disziplinen, für die die Analyse von Texten konstitutiv ist, eine immer stärker werdende Tendenz zur 'Linguistisierung' zu beobachten ist. So haben sich z.B. eine am linguistischen Strukturalismus orientierte Strukturalistische Stilistik (vgl. z.B. RIFFATERRE 1971) und eine Strukturalistische Poetik (vgl. STANKIEWICZ 1974) oder eine an der generativen Transformationsgrammatik orientierte Generative Poetik (vgl. z.B. VAN DIJK 1972a) konstituiert.⁸

1.3. Im Zusammenhang mit dieser allgemeinen Entwicklung muß auch die Tendenz gesehen werden, die Metrik auf der Basis der modernen Linguistik neu zu begründen.⁹

Anfänge dieser Entwicklung sind bereits relativ früh zu verzeichnen und gehen auf einige Arbeiten DE GROOTs und vor allem auf die zahlreichen Veröffentlichungen JAKOBSONs zurück. Erste Ansätze finden sich bereits bei DE GROOT (1918) und bei JAKOBSON (1923) und als bahnbrechend können JAKOBSON (1933) und DE GROOT (1946) angesehen werden. Aber auch neben den Veröffentlichungen JAKOBSONs und DE GROOTs gibt es vor allem in den osteuropäischen Ländern schon relativ früh eine Vielzahl von linguistisch orientierten Arbeiten zur Metrik. Der Hauptgrund für diese Entwicklung, die bis in die heutige Zeit anhält,¹⁰ dürfte in der engen Verbindung von linguistischen und literaturwissenschaftlichen Fragestellungen sowohl im russischen Formalismus als auch im russischen und tschechischen Strukturalismus zu sehen sein (vgl. EIMERMACHER 1969, ERLICH 1965).

In jüngerer Zeit hat die Linguistisierung der Metrik allgemein weiter zugenommen und der amerikanische Linguist J. LOTZ meint 1958 auf der interdisziplinären 'Conference on Style' zum Verhältnis von Metrik und Linguistik: "Since all metric phenomena are language phenomena, it follows that metrics is entirely within the competence of linguistics" (LOTZ 1964:137; Sperrung R.G.)

Einen wissenschaftsgeschichtlich bedeutsamen Einschnitt bildet das Jahr 1966. Während bis zu diesem Zeitpunkt vor allem Verfahren der strukturellen Phonologie auf die Metrik übertragen wurden - die Ansätze, zu denen auch noch der wichtige Beitrag von CHATMAN (1965) zu rechnen ist, werden deshalb auch als strukturalistische Metrik bezeichnet (vgl. z.B. LOTZ 1942, FOWLER 1966) - werden seit 1966 (vgl. HALLE/KEYSER 1966) auch zunehmend Verfahren der generativen Linguistik auf die Metrik übertragen. Diese Ansätze werden im Anschluß an BEAVER (1968a) als generative Metrik bezeichnet und im Augenblick kann wohl, zumindest was die Anzahl der Arbeiten zur generativen Metrik betrifft, von einem neuen Höhepunkt in der Entwicklung einer Theorie der Metrik gesprochen werden.¹¹

1.4. Neben und zum Teil auch in Verbindung mit der Linguistisierung der Metrik hat wiederum vor allem in Osteuropa eine Mathematisierung der Metrik stattgefunden. Der Schwerpunkt dieser Bewegung, die vor allem von Mathematikern und Linguisten ausging, liegt in der Anwendung statistischer und informationstheoretischer, d.h. quantitativer Methoden bei der Analyse formaler Textstrukturen. Die Verwendung dieser Verfahren in der Metrik darf dabei ebensowenig wie die Tendenz zur Linguistisierung isoliert gesehen werden, sondern muß im Zusammenhang mit einer allgemeinen Tendenz zur Mathematisierung betrachtet werden. Dieser kurze Hinweis auf die Möglichkeit der Anwendung mathematischer Verfahren mag an dieser Stelle vorerst genügen, da in den Kapiteln 11, 12 und 13 eine tiefergehende Diskussion dieses Ansatzes erfolgen soll.

1.5. Bisher habe ich drei Ansätze zur Neubegründung der Metrik unterschieden: die strukturalistische, die generative und die mathematische Metrik.¹² Neben diesen drei Ansätzen müssen noch die experimentalphonetischen und experimentalpsychologischen Arbeiten zur Metrik genannt werden, die im folgenden unter den Terminus 'experimentelle Metrik' subsumiert werden sollen. Diese Tradition, die bereits auf die kymographischen Studien von BRÜCKE (1871) zurückgeht und der z.B. die häufig zitierten Arbeiten von MEUMANN (1894), VERRIER (1909) und SCRIPTURE (1929) zuzurechnen sind, hat auch in jüngerer Zeit zahlreiche Arbeiten hervorgebracht.

1.6. Bei allen methodischen Divergenzen ist diesen vier Richtungen weitgehend gemeinsam das Streben nach Objektivierung von Deskriptions- und Analysemethoden und damit die Überwindung des vor allem in den hermeneutisch orientierten Arbeiten zur Metrik weit verbreiteten Subjektivismus.¹³ Da ich die methodologische Forderung nach Objektivität als grundlegend für eine empirische Wissenschaft wie die Metrik ansehe und da m.E. gerade von Seiten der strukturalistischen, generativen, mathematischen und experimentellen Metrik ein positiver Beitrag zur Weiterentwicklung einer Metriktheorie zu erwarten ist, werde ich - allerdings mit unterschiedlicher Gewichtung - vor allem diese Richtungen diskutieren.¹⁴

Dazu bedarf es jedoch zuerst einer Explikation sowohl des allgemeinen metatheoretischen Rahmens dieser Arbeit als auch einiger wichtiger wissenschaftstheoretischer Begriffe wie Objektivität, empirische Wissenschaft oder Theorie.

2. ZUR WISSENSCHAFTSTHEORETISCHEN BASIS EINER EMPIRISCHEN METRIK

2.1. Ich folge in dieser Arbeit dem wissenschaftstheoretischen Paradigma der Analytischen Wissenschaftstheorie, die neben Richtungen wie Hermeneutik und Dialektik als grundlegendes Paradigma der modernen Wissenschaftstheorie angesehen werden muß (vgl. RADNITZKY 1973).

Innerhalb der Analytischen Wissenschaftstheorie¹ stütze ich mich vor allem auf Ergebnisse des Kritischen Rationalismus, der auch bei den in letzter Zeit unternommenen Versuchen zur Neubegründung der Literaturwissenschaft eine zentrale Rolle spielt (vgl. z.B. PASTERNAK 1975, SCHMIDT 1975, EIBL 1976, GROEBEN 1976). Der Grund für diese Entscheidung ist darin zu sehen, daß m.E. eher von Seiten der Analytischen Wissenschaftstheorie und speziell des Kritischen Rationalismus als z.B. von Seiten der Hermeneutik ein Beitrag zur Lösung der in dieser Arbeit behandelten Probleme zu erwarten ist.

Im folgenden soll zuerst ein knapper Überblick über einige für die weitere Argumentation wichtige Grundgedanken des Kritischen Rationalismus erfolgen. Mit diesem Überblick soll gleichzeitig der Rahmen für einige später diskutierte Einzelprobleme wie Begriffsbildung, Erklärung oder Quantifizierung abgesteckt werden. Eine Auseinandersetzung mit aktuellen Gegenpositionen wie Kritische Theorie (vgl. z.B. ADORNO u.a. 1974) oder Konstruktivismus (vgl. z.B. HOLZKAMP 1972, ALBERT/KEUTH 1973) und eine ausführliche Darstellung der Kontroversen innerhalb des Kritischen Rationalismus, so z.B. des "epistemologischen Anarchismus" (vgl. FEYERABEND 1975), der Methodologie der Forschungsprogramme" (vgl. LAKATOS 1974, 1975) oder der Argumente von KUHN (1970, 1974), kann in diesem Kontext nicht geleistet werden.

2.2. Der Kritische Rationalismus, der vor allem auf die Arbeiten von POPPER zurückgeht (vgl. POPPER 1972a, 1972b, 1973) und in Deutschland besonders durch ALBERT (vgl. z.B. ALBERT 1969, 1971) vertreten wird, sieht das Ziel empirischer Wissenschaften

in der Konstruktion wahrer, informativer und präziser Theorien als Interpretation eines bestimmten Objektbereiches (vgl. auch KUTSCHERA 1972:391 ff). Gehen wir vom Sprachobligat wissenschaftlicher Erkenntnis aus (vgl. LEINFELLNER 1967:15), kann unter einer Theorie ein (deduktives) System von nomologischen Hypothesen (Gesetzen) verstanden werden.² Theorien sollen die Erklärung und Prognose konkreter raumzeitlich gegebener Tatbestände (singuläre Ereignisse) ermöglichen (vgl. POPPER 1957). Sie können als Netze aufgefaßt werden, die wir auswerfen, um die Welt einzufangen (vgl. POPPER 1973:31). Sie steuern unsere Wahrnehmung; es gibt daher keine theorieunabhängigen Beobachtungsdaten.³ Beobachtungen sind für POPPER "Interpretationen im Lichte von Theorien" (POPPER 1973:72). Theorien als Systeme von raum- und zeitlosen Gesetzen haben eine unendliche Extension und können aufgrund der ungelösten Induktionsproblematik (vgl. hierzu HOLZKAMP 1968:71-91) auch nicht durch Beobachtungsdaten begründet werden. Sie können nicht verifiziert, sondern lediglich falsifiziert werden, d.h. an der Realität scheitern.

Je öfter eine Theorie der kritischen Prüfung an der Realität standhält, desto höher ist ihr Bestätigungsgrad. Aus einer absolut sicheren Erkenntnisquelle gewonnene endgültig wahre Erkenntnisse gibt es für den Kritischen Rationalismus nicht (Fallibilismus). Es gibt lediglich eine mehr oder minder gute Annäherung an die Wahrheit, wobei Wahrheit hier im Sinne der Tarskischen Korrespondenztheorie (TARSKI 1935,1944) als Übereinstimmung von Aussagen mit der Realität aufzufassen ist. Neben der Übereinstimmung mit der Realität wird auch die logische Konsistenz einer Theorie als Kriterium für ihre Wahrheit angesehen.⁴

Grundlegend für den Kritischen Rationalismus ist die Forderung, daß einer Aussage dann und nur dann das Prädikat empirisch zugeschrieben werden soll, wenn sie falsifizierbar ist. Die Forderung nach empirischer Falsifizierbarkeit wird damit zum Abgrenzungskriterium zwischen empirischen und nicht-empirischen, z.B. metaphysischen oder mathematischen Aussagen und ersetzt das an das Prinzip der Verifikation gebundene empiristische Sinnkriterium, aufgrund dessen nicht verifizierbare Aussagen - wie z.B. metaphysische Aus-

sagen - als sinnlos anzusehen sind.⁵

Diese Fassung des Abgrenzungskriteriums führt jedoch zu problematischen Konsequenzen. Sowohl generelle Existenzhypothesen (aufgrund ihrer unendlichen Extension) als auch probabilistische Aussagen können nicht falsifiziert werden und sind deshalb innerhalb einer empirischen Wissenschaft nicht zugelassen.⁶

Diese Konsequenz halte ich vor allem in bezug auf probabilistische Aussagen für wenig sinnvoll. Es erscheint mir eher angebracht, lediglich zu verlangen, daß empirische Aussagen intersubjektiv an der Realität überprüfbar sein müssen, so z.B. probabilistische Aussagen mit Hilfe statistischer Verfahren. Zentral ist hierbei die Forderung nach Intersubjektivität, verstanden als intersubjektive Kritik im rationalen Wissenschaftsdiskurs. Denn erst die Intersubjektivität der Kritik ermöglicht die Objektivität wissenschaftlicher Sätze (vgl. POPPER 1973:18).

Gegenstand der Kritik ist lediglich die Begründung empirischer Aussagen. Die Frage ihrer Entstehung ist in diesem Zusammenhang irrelevant. Deshalb ist auch hinsichtlich der Geltung empirischer Aussagen strikt zwischen dem sog. Entstehungs- und dem sog. Begründungszusammenhang⁷ zu trennen. Gegenstand der Wissenschaftstheorie als normativer Wissenschaftslehre ist lediglich der Begründungszusammenhang, während für den Entstehungszusammenhang empirische Disziplinen wie Wissenssoziologie oder Wissenschaftsgeschichte als zuständig betrachtet werden.⁸

Neben den Problemen im Zusammenhang mit der Falsifikation probabilistischer Sätze hat vor allem das sog. Basisproblem (vgl. POPPER 1973:60ff) zur Kritik am POPPERschen Falsifikationsprinzip geführt. Die Überprüfung einer Theorie oder einer nomologischen Hypothese geschieht stets vermittelt (singulärer) Existenzaussagen. So ist z.B. eine nomologische Hypothese $\forall x[F(x) \rightarrow G(x)]$ dann falsifiziert, wenn die singuläre Existenzaussage $\exists x[F(x) \wedge \neg G(x)]$ als wahr akzeptiert wird. Solche singulären Existenzaussagen, die POPPER Basissätze nennt, sind jedoch sprachvermittelt und (alltags)theoriegeleitet. Sie haben ebenfalls lediglich hypothetischen Charakter und fallen ihrerseits wieder unter die Forderung

der intersubjektiven Nachprüfbarkeit.⁹ Letztendlich entscheidet bei POPPER deshalb der Konsens der Forscher darüber, ob ein Basisatz als Falsifikator akzeptiert wird (vgl. POPPER 1973:71ff).

Der verfeinerte (*sophisticated*) Falsifikationismus, der auf LAKATOS zurückgeht (vgl. z.B. LAKATOS 1974, 1975), revidiert deshalb die Poppersche Konzeption des Basisproblems.¹⁰ Auch die Basissätze werden als von Theorien generiert angesehen, so z.B. von sog. Beobachtungs- oder Instrumententheorien.¹¹ Sie entscheiden dann nicht mehr über die Falsifikation einer Theorie, sondern zeigen lediglich Inkonsistenzen zwischen mindestens zwei Theorien an (vgl. LAKATOS 1975). Eine Theorie wird nur dann eliminiert, wenn eine alternative Theorie mit einem höheren Informationsgehalt und einem höheren Bewährungsgrad existiert. Diese Fassung des Eliminationsprinzips führt dazu, daß der Poppersche monotheoretische Standpunkt zugunsten eines dynamischen Theorienpluralismus aufgegeben wird (vgl. auch SPINNER 1974).

2.3. Das in diesem Kapitel explizierte Verständnis einer empirischen Wissenschaft soll im folgenden bei der Betrachtung metrischer Texte als methodologische Norm zugrunde gelegt werden. Dabei gehe ich davon aus, daß die metatheoretischen Normen als topologische und nicht als klassifikatorische Begriffe konzipiert sind und daß Untersuchungen zur Metrik somit ein mehr oder minder großer Grad an 'Empirizität' zugeschrieben werden kann.

Es ist nun zu fragen, wie eine Metriktheorie zu konzipieren ist, die den explizierten metatheoretischen Normen möglichst weitgehend entspricht.

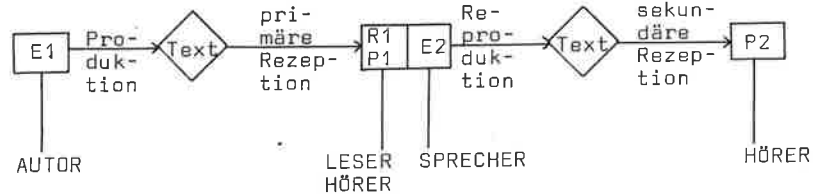
3. EIN RAHMENMODELL ZUR ANALYSE METRISCHER TEXTE

Die bisherige Forschung zur Metrik ist dadurch gekennzeichnet, daß eine Vielzahl verschiedener Aspekte metrischer Texte untersucht worden sind, ohne zumeist die Forschung an einem umfassenderen Rahmenmodell zu orientieren. Erst in jüngster Zeit sind vereinzelte Versuche unternommen worden, die Metrik als Teilkomponente eines umfassenderen Modells zu konzipieren (vgl. z.B. VAN DIJK 1972:210 ff, SIEVEKE 1973).

3.1. Betrachtet man sprachliche Texte als das Produkt komplexer Kommunikationsprozesse, die selbst wieder als Teile von komplexen Tätigkeitsakten anzusehen sind¹, bietet es sich an, metrische Texte im Rahmen einer kommunikationstheoretisch orientierten empirischen Textwissenschaft zu untersuchen. Vorschläge zur Konzipierung einer solchen Textwissenschaft sind aus verschiedener Richtung vor allem im Zusammenhang mit dem Versuch einer Neukonzeption der Literaturwissenschaft als Teildisziplin einer empirischen Kommunikationswissenschaft gemacht worden.²

Eine solche Textwissenschaft untersucht nicht mehr allein isolierte Texte, sondern die "Interaktion von Texten, Kommunikationsakten und Kommunikationssituationen" (SCHMIDT 1973:20). Konstitutiv für diesen Ansatz ist damit - neben der immer häufiger zu findenden expliziten Ausrichtung am Paradigma der Analytischen Wissenschaftstheorie - die pragmatische Dimension der Sprache.³

Aus der Sicht einer solchen kommunikationstheoretisch-textuell orientierten Linguistik/Literaturwissenschaft ist der Prozeß der Produktion und Rezeption metrischer Texte als ein spezieller Fall sprachlicher Kommunikation anzusehen, der somit anhand eines Kommunikationsmodells abgebildet werden kann.



E1 = primärer Expedient
 E2 = sekundärer Expedient
 R1 = primärer Rezipient
 P1 = primärer Perzipient
 P2 = sekundärer Perzipient

Abb. 3.1: textwissenschaftliche Kommunikationskette

Ein Textproduzent (Autor) produziert als primärer Expedient einen schriftlichen oder mündlichen Text. Der Text wird vom Leser bzw. Hörer, der als primärer Rezipient bzw. Perzipient fungiert, empfangen. Der Leser bzw. Hörer kann nun zum sekundären Expedienten werden, indem er als Sprecher den Text reproduziert und an einen Hörer sendet, der die Rolle eines sekundären Perzipienten einnimmt.

Anhand dieses einfachen Modells⁴ kann die Forschungssituation innerhalb der Metrik verdeutlicht werden. Kennzeichnend vor allem für literaturwissenschaftlich-hermeneutische Arbeiten zur Metrik ist die Tatsache, daß der Forscher (Metriker/Literaturwissenschaftler) gleichzeitig sowohl als Rezipient als auch als Analysator fungiert. Diese Verschmelzung der internen Rezipientenhaltung (Leser) mit der externen Beobachterposition (Forscher), bei der Rezeption und Interpretation zusammenfallen, führt jedoch wegen der ästhetisch bedingten Streubreite der Werkkonkretisationen zu einer hohen Subjektivität der Forschungsergebnisse.⁵ Intersubjektiv überprüfbare Resultate werden erst durch eine Beschränkung auf die externe Beobachterposition ermöglicht. Denn "die in einer Kommunikationskette

sich abspielenden Prozesse können nur von einem außerhalb der Kette stehenden externen Beobachter hinreichend exakt beschrieben werden, einem Beobachter, dem sämtliche Glieder einer Kette zugänglich sind". (MEYER-EPPLER 1969:5f)

Für eine empirisch-kommunikationswissenschaftliche Metrik ist deshalb eine strikte Trennung von Interpretation und Rezeption zu fordern. Nicht das eigene Rezeptionserlebnis, sondern die methodisch objektive Untersuchung des Verhaltens beliebiger Rezipienten zusammen mit einer intersubjektiv überprüfbaren Analyse der materialen Textstrukturen ergibt die Datenbasis für eine erfahrungswissenschaftliche Interpretation (vgl. GROEBEN 1976:126).

3.2. Um die verschiedenen Perspektiven bei der Analyse der Produktion und Rezeption metrischer Texte deutlicher charakterisieren zu können, soll anhand des Kommunikationsmodells zwischen verschiedenen Typen metrischer Kommunikation differenziert werden.

Berücksichtigt man den Übertragungskanal (optisch/akustisch), können vier Hauptkommunikationstypen unterschieden werden:

1. schriftliche Produktion → primäre visuelle Rezeption
2. mündliche Produktion → primäre auditive Rezeption
3. schriftliche Produktion → primäre visuelle Rezeption
→ mündliche Reproduktion (Rezitation)
→ sekundäre Rezeption
4. mündliche Produktion → primäre auditive Rezeption
→ mündliche Reproduktion → sekundäre Rezeption

Die Typen 2 und 4, die verschiedenen Formen mündlicher 'Dichtung' kennzeichnen, spielen u.a. aufgrund des immer stärker werdenden Einflusses des Mediums Schrift bei der Literaturproduktion in der Metrik lediglich eine marginale Rolle. Im Vordergrund des Forschungsinteresses steht vor allem der Typ 1 und zum Teil auch der Typ 3.

Je nachdem, welche Komponente des Kommunikationstyps schwerpunktmäßig untersucht wird, kann analytisch zwischen drei Forschungstypen differenziert werden:

1. autorzentrierte Produktionsanalyse
 - 1.1. schriftliche Produktion
 - 1.2. mündliche Produktion
2. textzentrierte Analyse
3. rezipientenzentrierte Analyse
 - 3.1. primäre visuelle Rezeption
 - 3.2. primäre auditive Rezeption
 - 3.3. Reproduktion (Rezitation)
 - 3.4. sekundäre Rezeption

Die Forschungstypen 1 und 3 sind der experimentellen Metrik zuzuordnen. Konstitutiv für diese Typen ist das Arbeiten mit Versuchspersonen.

Die textzentrierte Analyse ist charakteristisch für die strukturalistische, generative und mathematische Metrik. Sie abstrahiert vom Produktions- und Rezeptionsprozeß, versucht also nicht, die Konkretisation eines Werkes zu erfassen. Um dem Kriterium der Intersubjektivität zu genügen, muß sie sich vor allem auf den objektiv-materialen Aspekt eines Textes beschränken. Die in diesem Kontext angewandten Verfahren werden deshalb auch von GROEBEN (1972:175ff) "material-objektive Verfahren der Realitätsprüfung" genannt. Sie sind - vermutlich u.a. aus forschungspraktischen Gründen (keine Versuchspersonen) - die innerhalb der Metrik bisher hauptsächlich verwendeten Verfahren.

3.3. Für die Metrik ist vor allem die Analyse der Ausdrucksseite von Texten von Bedeutung. Denn bei der für die Metrizität von Texten konstitutiven Periodizität handelt es sich, sieht man einmal von der relativ seltenen visuellen Dichtung ab, vor allem um lautliche Periodizität.

Bei der Beschreibung der Ausdrucksseite metrischer Texte ist zwischen der phonetischen und der phonologischen Beschreibungsebene zu trennen. Die Forschungstypen 1 und 3 beschreiben die Ausdrucksseite von Texten auf der phonetischen Ebene. Auf dieser Ebene wird ein Text als konkretes Sprechereignis z.B. mit den Deskriptions-

methoden der artikulatorischen, akustischen und auditiven Phonetik untersucht. Das Resultat solcher Untersuchungen sind z.B. Aussagen über die Rezitation eines bestimmten Textes durch einen oder mehrere Sprecher oder über die Rezeption des rezitierten Textes durch einen oder mehrere Hörer.

Der lautlich aktualisierte Text ist jedoch nicht mit "d e m Text" gleichzusetzen.⁶

Will die Metrik zu allgemeinen, überindividuellen Aussagen über die lautliche S t r u k t u r von Texten gelangen, muß sie von der konkreten lautlichen Aktualisation eines Textes abstrahieren und den Text auf der phonologischen Ebene beschreiben. Die mit Hilfe der phonologischen Analyse gewonnenen Einheiten bilden dann die Grundlage für die weitere, z.B. statistisch orientierte Textanalyse.

Auf die Notwendigkeit der phonologischen Beschreibung hat schon MUKAŘOVSKÝ (1931:287) mit Nachdruck hingewiesen:

"L'application du point de vue phonologique est donc d'une utilité essentielle pour l'analyse du côté phonique de l'oeuvre d'art littéraire: Elle aide à établir une délimitation entre les qualités purement acoustiques (que l'on pourrait aussi appeler 'déclamatatoires'), facultatives au point de vue de l'oeuvre, et entre ces éléments phoniques qui sont des parties intégrantes de sa structure.

Das Abstraktionsniveau der phonologischen Beschreibung unterscheidet sich natürlich, je nachdem welche phonologische Theorie zugrunde gelegt wird. Aus der Sicht der generativen Phonologie ist noch eine abstraktere, zugrundeliegende Repräsentationsebene anzusetzen (vgl. MAYERHALER 1974:53ff), auf deren Bedeutung für die Metrik bzw. Poetik KIPARSKY (1968,1972) und SCHÄDLICH (1969) hingewiesen haben. KIPARSKY (1972:175) charakterisiert diese zusätzliche Ebene im Hinblick auf die Metrik folgendermaßen: "The phonological representations to which the metrical constraints apply are neither phonetic nor morphophonemic, but are intermediate representations..." SCHÄDLICH (1969:48f) meint im Hinblick auf die

Poetik: "Die Poetik hätte also nicht mit dem aktuellen Redesignal und nicht mit einer voll spezifizierten phonetischen Repräsentation, sondern mit einer höheren Abstraktionsstufe zu operieren, die zwischen der phonologischen und der voll spezifizierten phonetischen Repräsentationsebene der generativen Phonologie liegt."

Ebenso wie auf der phonologischen Beschreibungsebene können auch auf der phonetischen Beschreibungsebene je nach theoretischem Standpunkt verschiedene Abstraktionsebenen unterschieden werden. Der Unterschied "phonetisch vs. phonologisch" korrespondiert somit nur sehr bedingt mit dem Unterschied "konkret vs. abstrakt".

Die phonologische Orientierung ist charakteristisch für die textzentrierte Analyse der Ausdrucksseite metrischer Texte. Hierbei ergibt sich jedoch das Problem, daß bei schriftlichen Texten die vom Autor intendierte Lautstruktur anhand der Schrift rekonstruiert werden muß.⁷ Die Schrift repräsentiert allerdings die Lautstruktur nur partiell. Vor allem Phänomene wie Intonationsverläufe oder Pausenverteilung werden durch die Schrift lediglich unvollständig wiedergegeben. Da außerdem gerade bei metrischen Texten die lautliche Rekonstruktion zuweilen auch von der inhaltlichen Interpretation eines Textes abhängt, sollte sich der Forscher auch bei der phonologischen Analyse nicht allein auf seine eigene sprachliche bzw. metasprachliche Kompetenz verlassen, sondern durch Einbeziehung von Versuchspersonen die Objektivität der Analyse zu erhöhen suchen.

3.4. Die textzentrierte Analyse spielt aber auch im Hinblick auf die produktions- bzw. rezeptionszentrierte Analyse eine wichtige Rolle.

Im Normalfall kann der Prozeß der Produktion metrischer Texte nicht direkt untersucht werden, da nur bei zeitgenössischen Autoren eine direkte Datenerhebung z.B. mit Hilfe von Befragung oder teilnehmender Beobachtung möglich ist. Es kann somit in der Regel lediglich versucht werden, den Produktionsprozeß aufgrund von Textkorrelaten z.B. mit Hilfe biographischer, persönlichkeitspsychologischer oder literatursoziologischer Erklärungsmodelle zu rekonstruieren (vgl. GROEBEN 1972:27ff).

Auch bei der Rezeptionsanalyse kommt der Textanalyse eine wichtige Rolle zu. Denn auch sozialpsychologische oder experimentalphonetische Rezipientenanalysen können nicht von einer Analyse der materialen Textstrukturen absehen. Diese sind vielmehr als "Ermöglichungsgrund für Rezipientenverhalten (Konkretisationen etc.) auf jeden Fall in die empirische Untersuchung mit einzubeziehen" (GROEBEN 1976:127).

Der textzentrierten Analyse kommt auch aus diesem Grunde innerhalb der Metrik eine zentrale Bedeutung zu.

3.5. Beim Aufbau einer empirischen Theorie spielt die Begriffsbildung eine wichtige Rolle.

Im Anfangsstadium wissenschaftlicher Forschung entstammen die verwendeten Begriffe vorwiegend dem Vokabular der Alltagssprache. Begriffe der Alltagssprache sind jedoch in der Regel in dreifacher Hinsicht inexakt (vgl. ESSLER I 1970:56ff, OPP 1976:226ff):

- (a) Begriffe sind oft v a g e, d.h. die Regeln, nach denen sie verwendet werden, sind nicht immer eindeutig.
- (b) Begriffe sind m e h r d e u t i g (Homonymie)
- (c) Begriffe werden i n k o n s i s t e n t verwendet und zwar sowohl von verschiedenen Personen (interpersonale Inkonsistenz) als auch von ein und derselben Person (intrapersonale Inkonsistenz).

Wegen ihrer Inexaktheit sind Begriffe der Alltagssprache für die Formulierung intersubjektiv überprüfbarer Hypothesen wenig geeignet. Sie müssen deswegen entweder z.B. durch Definition oder Explikation normiert werden oder durch neu geschaffene Termini ersetzt werden. Die Exaktheit wissenschaftlicher Begriffe ist jedoch lediglich notwendiges, nicht aber hinreichendes Beurteilungskriterium.

Neben ihrer Inexaktheit sind Begriffe der Alltagssprache u.a. auch aufgrund ihres vorwiegend deskriptiven Charakters für die Theoriebildung nur begrenzt brauchbar. Die weitere Entwicklung einer wissenschaftlichen Disziplin führt deswegen stets zur Bildung von theoretischen Begriffen, die nur indirekt auf Beobachtungsdaten bezogen sind und zur Erklärung und Prognose von singulären Ereignissen benutzt werden können (vgl. HEMPEL 1974). Ein Beispiel für

ein solches theoretisches Konstrukt ist der Begriff der Kompetenz in der generativen Transformationsgrammatik.

Ein wichtiges Kriterium für die Beurteilung von Begriffen ist ihre theoretische Fruchtbarkeit (vgl. OPP 1976:234ff). Als theoretisch fruchtbar gilt ein Begriff dann, wenn er die Konstruktion brauchbarer empirischer Theorien ermöglicht. Diese Sicht der Funktion wissenschaftlicher Begriffe, die kennzeichnend für den kritischen Rationalismus ist, basiert erkenntnistheoretisch auf einem strikten Begriffsnominalismus⁸ und wendet sich gegen den in den Geistes- und Sozialwissenschaften weit verbreiteten Essentialismus, der die Funktion der Wissenschaft in der "Wesenserkenntnis" sieht (vgl. POPPER 1974:26-34).

Auch in nicht wenigen traditionellen Arbeiten zur Metrik findet sich eine solche essentialistische Deutung der Funktion wissenschaftlicher Begriffe. Außerdem sind viele der verwendeten Begriffe relativ inexakt.

Da auch die generative Metrik als jüngster Zweig der Metrik bisher keine befriedigende Metasprache entwickelt hat, ist IHWE (1975) zuzustimmen, der die Konstruktion einer genuinen Metasprache als dringendes Forschungsdesiderat betrachtet.⁹ Dieses Desiderat kann jedoch nur längerfristig realisiert werden, da die zu entwickelnde Metasprache anhand einer Vielzahl von Texten aus verschiedenen Sprachen und verschiedenen Epochen auf ihre Fruchtbarkeit hin zu überprüfen ist.

Im folgenden sollen nun zuerst einige für die Konzeption einer empirischen Metriktheorie m.E. grundlegende Begriffe geklärt werden.

4. TEXTE UND ZUFALLSFOLGEN

4.1. Eine zentrale Aufgabe der Metrik ist darin zu sehen, Kriterien für die Unterscheidung von metrischen und nicht-metrischen Texten anzugeben. Um eine solche Unterscheidung zu ermöglichen, ist jedoch zuerst der Begriff 'Text' genauer festzulegen.

Es existieren eine Vielzahl unterschiedlichster Textdefinitionen. Z.B. wird Text sprachimmanent als "ein durch ununterbrochene pronominale Verkettung konstituiertes Nacheinander sprachlicher Einheiten" (HARWEG 1968:148) oder handlungstheoretisch als "thematisch geordnete Menge von Instruktionen" (SCHMIDT 1973a: 273) definiert.

Für metrische Untersuchungen scheint mir folgende sehr allgemeine Definition fruchtbar zu sein:

Text =_{def} lineare Folge von sprachlichen Einheiten

In dieser Definition ist der Terminus 'Folge' im mathematischen Sinne zu verstehen, d.h. unter einer Folge wird eine Funktion f von der Menge der natürlichen Zahlen $\mathbb{N}$ in eine gegebene Menge M verstanden. Die Elemente von M werden auch als Grundsymbole bezeichnet. Eine Folge der Länge n sei eine Funktion von den ersten n natürlichen Zahlen ($\mathbb{N}_n$) in M .

Die Definition soll anhand des folgenden Beispiels verdeutlicht werden.

Es sei $M = \{a, b\}$, $n = 4$ und $f: \mathbb{N}_n \rightarrow M$.

Dann ist

$\{(1,a), (2,b), (3,a), (4,b)\} = abab$

eine mögliche Folge der Länge 4, d.h. es gilt

$abab \in f \subset \mathbb{N}_n \times M$.

Im Sinne dieser Definition kann unter einem Text z.B. eine Folge von Wörtern, eine Folge von Morphemen oder auch eine Folge von Silben verstanden werden.

4.2. Von besonderem Interesse für die Metrik ist die Betrachtung von sog. Binärfolgen, d.h. von Folgen, die nur aus den beiden Grundsymbolen 0 und 1 aufgebaut sind. Solche Folgen, die auch als 0-1-Folgen bezeichnet werden, sind von einer Reihe von Mathematikern untersucht worden. Sie spielen gerade in jüngster Zeit innerhalb der Wahrscheinlichkeitstheorie eine bedeutende Rolle (vgl. z.B. SCHNORR 1971).

Besonders wichtig für die Metrik ist die Verbindung von algorithmischen und stochastischen Ansätzen innerhalb der Theorie der Binärfolgen, wodurch m.E. eine gemeinsame Fundierung sowohl der generativen als auch der quantitativen Metrik ermöglicht wird. Aufgrund der Komplexität der Problematik werde ich mich im folgenden auf eine z.T. stark vereinfachende Einführung einiger elementarer Begriffe beschränken müssen.¹ Ich stütze mich dabei vor allem auf die Arbeiten von JACOBS (1969,1970) und auf FISCHER (1972), der m.W. als erster die Anwendungsmöglichkeiten der Theorie der Binärfolgen innerhalb der Textwissenschaft aufgezeigt hat.

4.2.1. Betrachten wir folgende Beispiele von Binärfolgen der Länge $n = 16$

- (1) 0 1 0 1 0 1 0 1 0 1 0 1 0 1
- (2) 0 1 1 0 1 0 0 1 1 0 0 1 0 1 1 0
- (3) 1 0 0 1 0 1 1 0 1 0 1 0 0 0 1 1

Die drei Folgen unterscheiden sich offensichtlich in Bezug auf ihre Regelmäßigkeit. Die Folge (1) hat die Periode 01 und ist absolut regelmäßig. Die Folge (2) ist unregelmäßiger als die Folge (1). Ihr liegt jedoch offensichtlich trotzdem ein Konstruktionsprinzip zugrunde. Der Folge (3) scheint zumindest auf den ersten Blick kein systematisches Konstruktionsprinzip zugrunde zu liegen. Sie scheint eher das Produkt des Zufalls zu sein.

Wenn wir einmal annehmen, daß es sich bei der Folge (3) um eine sog. Zufallsfolge handelt - solche Folgen können mit Hilfe eines physikalischen Zufallsgenerators (z.B. Münzwurf) erzeugt werden - stellt sich die Frage, wie Abweichungen von Zufallsfolgen bzw. Zufallsfolgen selbst mathematisch erfaßt werden können.

Hier ergeben sich allerdings einige Probleme. Wir können die drei Folgen z.B. aufgrund der Häufigkeit der auftretenden Symbole charakterisieren. Nach diesem Kriterium unterscheiden sich die drei Folgen jedoch überhaupt nicht. Sie haben die gleiche Anzahl von Elementen und, wenn wir den herkömmlichen Wahrscheinlichkeitsbegriff zugrunde legen, die gleiche Wahrscheinlichkeit, nämlich 2^{-16} .

Betrachten wir nun die Folge

(4) 0 1 0 0 0 0 0 0 1 0 0 0 0 0 1 0

Diese Folge würden wir zumindest intuitiv als nicht zufällig ansehen, da die relative Häufigkeit der Elemente 0 und 1 nicht gleich ist, d.h. es ist $p_0 \neq p_1 \neq 0,5$. FISCHER (1972:196) kennzeichnet deshalb zufällige Binärfolgen durch die notwendige Bedingung: "In einer zufälligen Binärfolge treten die Elemente 0 und 1 mit der gleichen relativen Häufigkeit auf."

Akzeptiert man diese Bedingung, sollte sie zumindest dahingehend präzisiert werden, daß $p = 0,5$ nur für $n \rightarrow \infty$ gilt, d.h. bei einer zufälligen Binärfolge $X = x_1 x_2 \dots$ gilt nach dem starken Gesetz der großen Zahlen (vgl. JACOBS 1970:164):

$$\lim_{n \rightarrow \infty} \frac{1}{n} (x_1 + x_2 + \dots + x_n) = 0,5$$

Bei einer konkret vorliegenden endlichen Folge würden wir lediglich erwarten, daß die Anzahl der Einsen und Nullen ungefähr $n/2$ beträgt.

Die o.a. Bedingung führt jedoch zu einer m.E. für die statistische Textanalyse wenig wünschenswerten Einengung des Begriffs der Zufallsfolge auf den symmetrischen Fall, bei dem $p_1 = p_0 = 0,5$ ist. Es soll deshalb im folgenden der Begriff der Zufallsfolge auch auf asymmetrische Fälle, d.h. Fälle mit $p_1 \neq p_0 \neq 0,5$ ausgedehnt werden. Solche Folgen, die sich ebenfalls durch einen physikalischen Zufallsgenerator erzeugen lassen, werden von MARTIN-LÖF (1966:612ff) als "random sequences with respect to an arbitrary computable probability distribution" charakterisiert.

4.2.2. Um den Zufälligkeitsgrad einer Folge bestimmen zu können, muß jedoch, wie die Folgen (1) bis (3) zeigen, neben der Häufigkeitsverteilung der Elemente auch noch die Anordnung der Elemente innerhalb der Folge berücksichtigt werden.

Dies kann mit Hilfe verschiedener Methoden geschehen. Mit algebraischen und topologischen Methoden läßt sich z.B. die Zusammenhangs- und Repetitionsstruktur von Folgen erfassen (vgl. z.B. FISCHER 1969; 1970; 1970a). Mit Hilfe stochastischer Modelle kann z.B. die Wahrscheinlichkeit für das Auftreten bestimmter Periodizitäten bestimmt werden (vgl. Kap. 12).

Für die algorithmische Charakterisierung von Binärfolgen ist es zweckmäßig, den Aufbau von Folgen aus sog. Blöcken zu untersuchen (vgl. hierzu JACOBS 1969). Im Rahmen dieser Betrachtungsweise wird eine endliche Folge $A = a_1 \dots a_n$ von Nullen und Einsen auch als (0-1-) Block der Länge $|A| = n$ bezeichnet. $a_1, \dots, a_n$ heißt die 1-te, ..., n-te Komponente des Blocks. Mit Blöcken können verschiedene algebraische Operationen durchgeführt werden, wie z.B. Addition, Multiplikation und Spiegelung.

Unter einer Block-Addition² versteht man das Nebeneinanderstellen von Blöcken. Sind z.B. $A = a_1 \dots a_n$, $B = b_1 \dots b_m$ zwei 0-1-Blöcke, definiert man

$$AB = a_1 \dots a_n b_1 \dots b_m.$$

Für die Blocklänge gilt

$$|AB| = |A| + |B|.$$

Die Blockaddition ist assoziativ, jedoch nicht kommutativ.

Für die Spiegelung, die durch die Exponenten 1 bzw. 0 symbolisiert wird, gilt

$$0^1 = 1, 1^1 = 0, 0^0 = 0, 1^0 = 1, \text{ sowie } A^1 = a_1^1 \dots a_n^1 \text{ und } A^0 = A.$$

Man erhält z.B. $(01)^1 = 10$ oder $(1001)^1 = 0110$.

Die Multiplikation ist folgendermaßen definiert:

Sind $A = a_1 \dots a_n$, $B = b_1 \dots b_m$ 0-1-Blöcke, ist

$$A \times B = A^{b_1} \dots A^{b_m}$$

das Produkt von A und B. Die Länge des Produkts beträgt

$$|A \times B| = |A| \cdot |B|.$$

Man erhält z.B. $(01) \times (01) = 0110$ oder $(0110) \times (01) = 01101001$.

Mit Hilfe der Blockmultiplikation kann jetzt z.B. die Folge (2) - es handelt sich übrigens um eine sog. Morse-Folge - als Produkt $(01) \times (01) \times (01) \times \dots$ konstruiert werden.

Bereits die Operationen Addition, Multiplikation und Spiegelung erlauben es, eine Vielzahl verschiedener Binärfolgen algorithmisch zu erzeugen, bzw. vorliegende Folgen in ihrer Struktur zu beschreiben (vgl. JACOBS 1969).

Auf dem hier angedeuteten Weg ist es auch möglich, den Zufälligkeitsgrad von Folgen durch die Komplexität ihrer Beschreibung zu definieren. Der Zufälligkeitsgrad einer Folge ist um so höher, je komplexer der ihr zugrundeliegende Algorithmus ist. Für ideale Zufallsfolgen gilt, daß diese nicht algorithmisch berechenbar sind (vgl. JACOBS 1970:165).

4.3. Folgen aus sprachlichen Einheiten können nun daraufhin untersucht werden, inwieweit sie zufälligen Binärfolgen entsprechen.

Betrachten wir z.B. einen beliebigen deutschen Nachrichtentext, so zeigt dieser u.a. aufgrund der deutschen Akzentgesetze gewisse Regelmäßigkeiten in der Abfolge von betonten und unbetonten Silben.³

Analysiert man jedoch die Abfolge von betonten und unbetonten Silben in einem Text wie

Du liebes Kind, komm, geh mit mir!

(GOETHE, Erlkönig)

zeigt sich eine für das Deutsche ungewöhnliche absolute Periodizität.

Bezeichnen wir eine unbetonte Silbe mit 0 und eine betonte Silbe mit 1, entspricht dem Erbkönig-Vers die Folge 0 1 0 1 0 1 0 1. Diese könnte z.B. durch die Zuordnungsvorschrift

$$x \rightarrow \begin{cases} 0, & \text{falls } x \in 2\mathbb{N} + 1 \wedge x \leq 8 \\ 1, & \text{falls } x \in 2\mathbb{N} \wedge x \leq 8 \end{cases}$$

erzeugt werden. Die Folge hat die Periode 2, weil für den Wertebereich der Funktion f gilt:

$$f(x) = f(x + 2), \quad \text{für } x \leq 8$$

Blockalgebraisch könnte diese Folge auch als Addition von 4 0-1-Blöcken interpretiert werden.

Texte, die eine solche im Vergleich zur Gesamtsprache erhöhte numerische Regularität⁴ in der Silbenfolge aufweisen, sind von mir in Kap. 1 als metrische Texte bezeichnet worden. Die zusätzliche Regularität ist aber lediglich eine notwendige, jedoch keine hinreichende Bedingung für die Differenzierung zwischen metrischen und nicht-metrischen Texten. So wird z.B. eine Folge von 100 Silben, bei der lediglich die Zuordnung für alle einer Quadratzahl entsprechenden Positionen geregelt ist, üblicherweise als nicht-metrisch angesehen.⁵

Damit stellt sich die Frage, welche weiteren Bedingungen erfüllt sein müssen, damit objektiv beschreibbarer sprachlicher Periodizität das Prädikat metrisch zugeschrieben werden kann. Eine Klärung dieser Frage soll nun im Zusammenhang mit der Diskussion der Begriffe Rhythmus und Metrum erfolgen.

5. RHYTHMUS UND METRUM

Für die Metrik sind die Termini Metrum und Rhythmus von grundlegender Bedeutung. Diese Termini werden jedoch z.T. inkonsistent und z.T. auch unpräzise verwendet.¹ Deshalb sollen zuerst diese beiden Termini begrifflich abgegrenzt werden.

5.1. Der Rhythmus

Das Phänomen des Rhythmus ist vor allem von Psychologen, und zwar vorwiegend als wahrnehmungspsychologisches Problem, empirisch untersucht worden. Daneben haben zum speziellen Problem des Sprachrhythmus auch Phonetiker und Metriker empirische Untersuchungen durchgeführt.²

5.1.1. Es ist die grundsätzliche Unterscheidung zwischen temporalem und spatiolem Rhythmus³ gemacht worden (vgl. CHATMAN 1965:18), die in erster Annäherung als zeitliche bzw. räumliche Periodizität definiert werden können.

Ein Beispiel für spatialen Rhythmus ist eine nach einem Ordnungsprinzip angepflanzte Reihe von Bäumen. Meist wird jedoch ein solches Phänomen, sofern es gleichzeitig perzipiert wird, als Symmetrie bezeichnet (CHATMAN 1965:18). Beispiele für die Verwendung des Terminus Rhythmus im Sinne von zeitlicher Periodizität sind Begriffe wie Schlafrythmus, Stoffwechselrhythmus oder Rhythmus der Jahreszeiten.

Bei temporalem Rhythmus wird das Zeitkontinuum in Phasen gegliedert. Nach CHATMAN (1965:19) geschieht diese Gliederung mit Hilfe von Intervallen und Ereignissen ("events"). CHATMAN spricht - möglicherweise in Anlehnung an die Verwendungsweise der Termini *event* bzw. Ereignis in der Wahrscheinlichkeitstheorie - mit Absicht sehr allgemein von Ereignissen. Denn die Ereignisse können sehr unterschiedlicher Natur sein, da rhythmische Strukturen mit verschiedenen Sinnesorganen erfaßt werden können (CHATMAN 1965:24), so z.B. rhythmische Bewegung mit den Augen. In den meisten Fällen in denen von Rhythmus gesprochen wird, sind die Ereignisse jedoch Töne.

CHATMAN (1965:19) scheint Intervalle als notwendige Konstituenten rhythmischer Strukturen anzusehen. Dies ist m.E. jedoch weder für den außersprachlichen noch für den sprachlichen Bereich zutreffend. Wie Sonagramme von Gedichtrezitationen zeigen, finden sich häufig mehrere Silben hintereinander keine Pausen. Auch in der Musik ist es möglich, Rhythmus zu erzeugen, indem das lautliche Kontinuum in einzelne lautliche Ereignisse unterteilt wird, ohne daß zwischen den Ereignissen Intervalle liegen. Dies kann z.B. dadurch geschehen, daß die Frequenz oder die Amplitude von Teilen des Lautkontinuums systematisch variiert wird, so daß eine rhythmische Folge von Ereignissen perzipiert wird.⁴

CHATMAN (1965:20) unterscheidet im Anschluß an MC DOUGALL (1902) noch primären und sekundären Rhythmus.

Bei gleichen Ereignissen und gleichen Intervallen handelt es sich um primären Rhythmus, bei einer Gruppierung aufgrund von ungleichen Intervallen und/oder Ereignissen um sekundären Rhythmus. Ein Beispiel für primären Rhythmus wäre eine Folge wie - - - -, für sekundären Rhythmus — - — - — - oder - - - - - oder

Der sekundäre Rhythmus ist weit häufiger als der primäre Rhythmus. Sehr oft wird primärer Rhythmus auch als sekundärer Rhythmus perzipiert (vgl. JACOB 1918:95f; FRAISSE 1956:9f). So perzipieren wir z.B. den primären Rhythmus eines Metronoms nach kurzem Zuhören subjektiv als sekundären Rhythmus. Der Grund für diese automatische 'subjektive Rhythmisierung' (WUNDT) dürfte in der allgemeinen Funktionsweise menschlicher kognitiver Prozesse zu sehen sein, die notwendigerweise zu einer Strukturierung wahrgenommener Signalfolgen führt (vgl. BRUNER u.a. 1956)⁵.

Konstitutiv für den Rhythmus ist die Wiederkehr von gleichen Ereignissen. Gleichheit bedeutet aber nicht absolute physikalische Gleichheit, sondern die Ereignisse müssen lediglich als gleich empfunden werden. Denn sowohl Ereignisse als auch Intervalle werden, wie schon lange experimentell nachgewiesen ist, mit einer gewissen Variationsbreite⁶ realisiert, ohne daß der Eindruck von Gleichheit verloren geht: "The primary requisite for the unit groups is that they shall be alike, not that they shall be equal." (STETSON 1903:462).⁷ Allerdings erhöht

eine absolute Periodizität, wie WALLIN (1901:108) festgestellt hat, die Qualität der rhythmischen Empfindung:

"Absolute periodic or regular occurrences are not essential to the appreciation of rhythm, although absolute regularity improves the quality of the rhythmic impression. To engender a feeling of rhythm always requires a certain amount of periodicity, but the margin of irregularity is quite considerable."

Psychologisch gesehen handelt es sich um einen Kategorisierungsprozeß, bei dem Einheiten, die hinsichtlich ihrer Funktion als äquivalent angesehen werden, der gleichen Kategorie zugeordnet werden.⁸ Dieser Vorgang, der auch als funktionale Äquivalenzkategorisierung bezeichnet wird, führt zu der für die Wahrnehmung rhythmischer und metrischer Strukturen notwendigen Invariantenbildung (vgl. BRUNER u.a. 1956:2ff). Es handelt sich dabei um einen Sonderfall der jeglicher Art von Strukturerkennung zugrundeliegenden (perzeptorischen) Abstraktion (vgl. z.B. CORCORAN 1971). Welche mentalen Prozesse bei diesem Abstraktionsprozeß im einzelnen involviert sind, ist bisher nur z.T. geklärt. Für die Metrik ist u.a. von Bedeutung, daß auch die Wahrscheinlichkeitsstruktur der betreffenden Ereignisfolge eine Rolle spielt (vgl. BRUNER u.a. 1956:182ff).

Es dürfte deutlich geworden sein, daß bei der Definition von Rhythmus sowohl die objektiven Eigenschaften von Ereignisfolgen als auch das Perzeptionsverhalten von Versuchspersonen von Bedeutung ist.

Begriffen wie Schlafrhythmus oder Stoffwechselrhythmus liegt ein weites Verständnis von Rhythmus zugrunde, bei dem die objektiven Eigenschaften des betreffenden Ereignisses im Vordergrund stehen. Psychologen beziehen dagegen bei der Definition von Rhythmus in der Regel den Perzipienten mit ein und betonen die Interdependenz von Stimuluseigenschaften und Perzipientenverhalten.⁹ Ein deutliches Beispiel für eine solche Rhythmusdefinition findet sich bei WOODROW (1951:1232f):

"By rhythm, in the psychological sense, is meant the perception of a series of stimuli as a series of groups of stimuli. The successive groups are ordinarily of similar pattern and experienced as repetitive. (...) Objective characteristics of the stimulus series by no means completely determine the grouping."

5.1.2. Wird der Begriff Rhythmus auf die Sprache bezogen, sind die periodischen Elemente stets Silben bzw. Silbenfolgen. Auch der auf Sprache bezogene Rhythmusbegriff wird sowohl im perzipientenunabhängigen als auch im perzipientenabhängigen Sinne verwendet. Ein Beispiel für die perzipientenunabhängige Verwendung ist die folgende Äußerung von MALMBERG (1971:21):

"Jede gesprochene Äußerung ist rhythmisch. Schon die für die Bildung von Sprachlauten erforderliche Ausatmung setzt eine ziemlich regelmäßige Aktivität der Rippenmuskeln voraus."

Im Gegensatz zu dieser allgemeinen Verwendungsweise des Begriffes Rhythmus¹⁰ liegt jedoch z.B. dem Begriff des Prosarhythmus in der Regel eine perzipientenabhängige Rhythmusdefinition zugrunde.

In einer Vielzahl von Arbeiten zur Metrik wird jedoch kaum deutlich, welche der beiden Verwendungsweisen des Terminus Rhythmus zugrundegelegt wird. Um mögliche Mißverständnisse zu vermeiden, werde ich deshalb im folgenden den Rhythmusbegriff entsprechend der o.a. Definition von WOODROW, d.h. lediglich im perzipientenabhängigen Sinne verwenden und anstelle des perzipientenunabhängigen Rhythmusbegriffs Termini wie Periodizität oder Regularität gebrauchen.¹¹

5.2. Das Metrum

5.2.1. Viele Metriker sind der Meinung, daß sich Rhythmus und Metrum wie Genus und Spezies verhalten,¹² d.h., daß das Metrum lediglich als eine Sonderform des Sprachrhythmus anzusehen ist (vgl. hierzu CHATMAN 1965:12).

Andere Metriker betrachten jedoch das Metrum eher als Abstraktion.¹³ So verstehen z.B. PAUL/GLIER (1964:17) unter Metrum "das aus der Sprachgestalt zu *a b s t r a h i e r e n d e* ... Schema", das dem Rhythmus als der aktuellen und individuellen sprachlichen Realisierung des Metrums durch den Dichter zugrunde liegt.¹⁴ Hinsichtlich des Grades der Abstraktheit des Metrums herrscht allerdings keineswegs Einigkeit. TARLINSKAJA/TETERINA (1974:63ff) z.B. unterscheiden fünf verschiedene Abstraktionsebenen bei der Analyse metrischer Texte. Eine relativ konkrete Auffassung von Metrum findet sich vor allem bei statistisch arbeitenden Metrikern, die den Terminus Metrum eher im Sinne einer deskriptiv verstandenen durchschnittlichen Realisationsnorm (Erwartungswert) verwenden (vgl. z.B. BAILEY 1973). Einen hohen Grad von Abstraktheit weist dagegen die Definition von Metrum in der generativen Metrik auf, wo z.B. von HALLE/KEYSER (1971:140) unter Metrum "the encoding of a simple abstract pattern into a sequence of words" verstanden wird.

5.2.2. Ich werde im folgenden ebenfalls eine relativ abstrakte und sehr allgemeine Definition von Metrum zugrunde legen, die als tentative Arbeitsdefinition für die weitere Argumentation dienen soll und die sowohl einzelsprachlich als auch textsortenspezifisch zu konkretisieren ist:

Metrum =_{def} funktionales Konstruktionsprinzip einer rhythmischen Silbenfolge

Diese Definition berücksichtigt die objektiven Textcharakteristiken, den Textproduzenten und den Textrezipienten.

Durch die Einschränkung des Konstruktionsprinzips auf rhythmische Silbenfolgen wird vermieden, daß jegliche Art von zusätzlicher numerischer Regularität¹⁵ in der Abfolge von Silben als metrisch

klassifiziert wird (vgl. das Beispiel auf S. 22; vgl. auch BERNHART 1974).

Durch die Kennzeichnung des Konstruktionsprinzips als funktional (vgl. auch Kap. 10) soll angedeutet werden, daß dieses Prinzip nicht als sprachimmanente Gesetzmäßigkeit, sondern eher als eine der numerischen Regulierung von Silbenfolgen zugrundeliegende (literarisch-ästhetische) Konvention anzusehen ist.¹⁶

Das Metrum als Konvention ist häufig historisch tradiert und, wie z.B. in den normativen Poetiken des 16. und 17. Jahrhunderts, explizit als Vorschrift formuliert. Die Konvention kann aber auch, wie z.B. im Fall der serbokroatischen Volksepen, dem Produzenten/Rezipienten weitgehend unbewußt sein.

5.2.3. Als metrisch sind somit solche Texte anzusehen, die erstens eine rhythmische Struktur aufweisen und denen zweitens ein Metrum zugrunde liegt. Nicht-metrische Texte sollen jetzt als Prosa(texte) bezeichnet werden. Der Begriff Prosa(text) ist damit lediglich formal gekennzeichnet und wird nicht etwa, wie es häufig der Fall ist, als inhaltlicher Gegenbegriff zu 'poetischer Text' verstanden.

Benutzt man zur Charakterisierung der Opposition zwischen metrischen Texten und Prosatexten den Begriff der 'Markiertheit',¹⁷ sind metrische Texte, da sie weniger frequent und weniger 'natürlich' als Prosatexte sind, als markiert, Prosatexte dagegen als unmarkiert anzusehen (vgl. VALESIO 1971:50; LOTZ 1974:964).

Die vorgeschlagene Explikation des Prädikats 'metrisch' ist nicht mit der Explikation des Prädikats '*metrical*' in der generativen Metrik zu verwechseln.

Erstens betrachte ich 'metrisch' als Prädikat von Texten relativ zu einer bestimmten Sprache L und nicht, wie '*metrical*' meist in der generativen Metrik verstanden wird, als Prädikat von Versen (Texten) relativ zu einem bestimmten metrischen System (vgl. z.B. HALLE/KEYSER 1971:139ff). Zweitens habe ich die pragmatische Textdimension (Produktion/Rezeption) mit einbezogen (vgl. hierzu IHWE 1975:382,386).

'Metrisch' ist somit als dreistelliges Prädikat aufzufassen: M (Text, Produzent/Rezipient, Sprache L).¹⁸

Aus der Sicht der Binärfolgen ist das Metrum als Algorithmus zu interpretieren, der eine systematische Abweichung von einer Zufallsfolge erzeugt. Da der Zufälligkeitsgrad einer Folge u.a. anhand der Komplexität des zugrundeliegenden Algorithmus gemessen werden kann (vgl. Kap. 4), ist es auf diese Weise möglich, das Prädikat 'metrisch' auf der Textbeschreibungsebene zu quantifizieren. Dies bedeutet, daß Grade von Metrizität unterschieden werden können und auch Übergangsformen wie z.B. rhythmische Prosa oder freie Rhythmen¹⁹ anhand ihres Zufälligkeitsgrades objektiv charakterisiert werden können. Auf weitere Möglichkeiten der Analyse mit Hilfe der Theorie der Binärfolgen habe ich bereits hingewiesen.

5.3. Der Vers

Meist ist ein metrischer Text aus Teilfolgen (Blöcken) mit gleichem oder auch, wie z.B. im Fall der lyrischen Strophen des HORAZ, mit unterschiedlichem Metrum aufgebaut. In der Regel sind die Teilfolgen *typographisch* durch Anordnung in Zeilen und/oder *sprachlich* durch Grenzsignale (vgl. VALESIO 1971:61) wie syntaktische Einschnitte bzw. Pausen, Reim oder *syllaba anceps* im Griechischen und Lateinischen voneinander abgegrenzt. Teilfolgen dieser Art sollen als Verse, aus Versen aufgebaute Texte als Verstexte bezeichnet werden. Der Terminus 'Zeile', der in der Metrik zuweilen mit Vers gleichgesetzt wird, soll dagegen lediglich als typographische Einheit verstanden werden.²⁰ Verstexte sind entsprechend dieser Definition lediglich eine - allerdings sehr umfangreiche - Teilmenge metrischer Texte.

Der Vers ist ebenso wie das Metrum als (literarisch-ästhetische) Konvention anzusehen.

Meist sind Verse sowohl durch typographische als auch durch sprachliche Signale gekennzeichnet. Es finden sich allerdings auch metrische Texte mit fortlaufend geschriebenen, sprachlich markierten Versen. So hat nach HRABÁK (1961:245) M. GOR'KIJ einige Verse fortlaufend geschrieben. Allerdings handelt es sich dabei um Verse mit

traditionellen Zügen, die ohne Schwierigkeit als Vers erkannt werden. Ein weiteres Beispiel für fortlaufend geschriebene Verse sind die "Ballades françaises" von P. FORT. Außerdem besteht auch, wie BEAVER (1968a:155) bemerkt, die Möglichkeit, daß typographische und sprachliche Einteilung im Widerspruch stehen: "...real line length does not necessarily correspond to the line length as the poet arranges his verse v i s u a l l y." (Sperrung R.G.)

Als n o t w e n d i g e Konstituente des Verses ist somit lediglich das Metrum anzusehen. Die typographische Gestaltung spielt vor allem bei solchen Versen eine Rolle, die aufgrund ihrer sprachlichen Konstituenten kaum noch als Verse zu erkennen sind, wie dies z.B. beim *vers libre* der Fall ist.²¹

"Tant que dans une forme versifiée il y a moins d'éléments qui distinguent le vers de la prose il faut marquer d'autant plus clairement qu'il s'agit du vers et non de la prose...C'est pourquoi le vers libre exige un certain aspect graphique pour être perçu comme une forme du langage 'en vers'." (HRABÁK 1961:245)

Häufig wird in der Metrik die Einteilung eines Textes in Verse nicht weiter problematisiert. Dies ist auch bei vielen Vertretern der generativen Metrik der Fall²², so z.B. bei HALLE/KEYSER (1971), die offensichtlich von einer nicht hinterfragten typographischen Verseinteilung ausgehen (vgl. die Kritik bei KLEIN 1974: 44 und bei IHWE 1975:370; vgl. auch FOWLER 1976:27).

Eine Metriktheorie darf jedoch m.E. nicht einen rein typographisch definierten Versbegriff zugrundelegen, zumal die typographische Zeileneinteilung im Laufe der Zeit verändert worden sein kann. Sie muß auch eine Antwort darauf geben können, aufgrund von welchen sprachlichen Merkmalen ein gegebener metrischer Text als Verstext identifiziert werden kann und ob typographische und sprachliche Merkmale korrespondieren.²³

5.4. Metrum, rhythmische Struktur und rhythmische Realisation

5.4.1. Das Metrum ist als theoretisches Konstrukt selbst nicht direkt beobachtbar, sondern kann lediglich indirekt anhand seiner sprachlichen Realisation in Form des Rhythmus erfaßt werden.²⁴

Bei der Beschreibung der objektiven Charakteristiken des Rhythmus sind zwei Abstraktionsstufen zu unterscheiden: Zum einen die phonologische Beschreibung der rhythmischen S t r u k t u r eines Textes, zum anderen die phonetische Beschreibung des durch die rhythmische Struktur nur partiell determinierten, lautlich realisierten Rhythmus z.B. bei der Rezitation eines Textes.²⁵

Benutzt man ein von BIERWISCH (1966:81f) für das Verhältnis von Sprache und Sprechen verwendetes Bild, kann man das Verhältnis von rhythmischer Struktur und lautlich realisiertem Rhythmus, im folgenden 'rhythmische Realisation' genannt, "in Analogie zur Partitur einer Symphonie und ihren vielen möglichen Aufführungen verstehen, die von der kompositorisch fixierten Struktur determiniert, aber nicht mit ihr identisch ist: Jede Aufführung hat eine eigene akustische Existenzform, sie kann mehr oder weniger von der Partitur abweichen, sie enthält Varianten und Fehler und stellt eine spezielle Interpretation der Partitur dar." (vgl. auch HALLE/KEYSER 1966:191)

Die rhythmische Realisation läßt sich zumindest approximativ mit Hilfe der in Amerika von den Anglisten der Universitäten noch in den sechziger Jahren fast allgemein akzeptierten musikalischen Notation (vgl. WELLEK/WARREN 1966:145) beschreiben. Die musikalische Notation beschreibt jedoch nicht, wie häufig angenommen wird, das Metrum, da sie nicht von der individuellen lautlichen Variation und von der individuellen Textgestaltung abstrahiert (vgl. WIMSATT/BEARDSLEY 1964:195).

Die m.E. wichtige Unterscheidung von Metrum, rhythmischer Struktur und rhythmischer Realisation wird in einer nicht unbeträchtlichen Anzahl von Arbeiten zur Metrik wenig deutlich. Aber auch wenn die Notwendigkeit der Unterscheidung von drei Abstraktionsstufen gesehen wird, wird terminologisch in der Regel lediglich zwischen Metrum - zuweilen wird in der deutschsprachigen Literatur auch anstelle von Metrum der Terminus 'Versmaß' verwendet - und (Vers)rhythmus differenziert.²⁶ Daneben finden sich auch noch Dichotomien wie z.B. "estructura" vs. "realización" (RUIPÉREZ 1955:91) "meter" vs. "actualisation" (HALLE 1970:64) oder "invariant" vs. "variantes" (GOLDENBERG 1976:98).

Eine klare terminologische Trennung findet sich dagegen bei JAKOBSON (1964:364f), der die Ebenen des "verse design" (Metrum), "verse instance" (rhythmische Struktur) und "delivery instance" (rhythmische Realisation) unterscheidet.

5.4.2. Ebenso wie die rhythmische Struktur die rhythmische Realisation eines metrischen Textes nur partiell determiniert, determiniert auch das Metrum in der Regel nur partiell die rhythmische Struktur. Die meisten metrischen Texte weisen allein schon aus diesem Grund eine mehr oder minder große rhythmische Variation auf.

Eine weitere Quelle der rhythmischen Variation ist die Abweichung des Textproduzenten (Dichter) vom zugrundegelegten Metrum beim Streben nach individueller metrischer Gestaltung.²⁷ Außerdem kann auch durch die spezifische stilistische Textgestaltung rhythmisch differenziert werden. Die Mittel hierfür sind z.T. sprachspezifisch. So kann z.B. in tschechischen Gedichten mit einem durch die Tradition konventionalisierten absolut regelmäßigen Metrum durch die Verteilung der Wortgrenzen im Vers auditiv ein Eindruck von weitgehender Unregelmäßigkeit entstehen (vgl. MUKAŘOVSKÝ 1931:286f).²⁸

5.4.3. Der Rezipient metrischer Texte abstrahiert das Metrum z.B. aufgrund der Akzentverteilung, der Pausenverteilung oder auch aufgrund der typographischen Anordnung aus einem konkret vorliegenden Text. Dabei spielt vor allem bei komplizierten Metren die vorhergehende, meist in einem literarischen Lernprozeß erworbene Kenntnis des Metrums eine wichtige Rolle. Hat der Rezipient das Metrum internalisiert, kann das Metrum als psychisches Konzept²⁹ den weiteren Rezeptionsprozeß steuern (vgl. THOMPSON 1961:170; BAILEY 1975:22). Das kann sogar dazu führen, daß metrisch stark abweichende Verse als metrisch regelmäßig empfunden werden (vgl. LIGHTFOOT 1974:245).

Dieser Vorgang kann phänomenal als Ausgliederung einer Gestalt betrachtet werden (vgl. STETSON 1903:459). Eine solche 'metrische

Gestalt' genügt den EHRENFELSSchen Gestaltqualitäten der Übersummativität und der Transponierbarkeit. Sie wird als abgegrenzte, gegliederte, ganzheitliche Verlaufsgestalt perzipiert, deren Struktur trotz z.B. unterschiedlicher Rezitation (Transponierung) erhalten bleibt.

Unterschiedliche Rezitationen sind jedoch nicht in jedem Fall gestalterhaltend. Variiert z.B. die Grundfrequenz der Stimme bei verschiedenen Rezitationen, dann kann man hierbei von Transponierung im Sinne der Gestalttheorie sprechen. Dies gilt jedoch nicht, wenn z.B. bei einer Rezitation die betonte Silbe viermal so laut wie die unbetonte Silbe, bei einer anderen Rezitation jedoch nur zweimal so laut wie die unbetonte Silbe gesprochen wird. Obwohl das Metrum nicht verändert wird, findet bei einer solchen Veränderung der relativen Betonungsstärke eine Neuzentrierung der Gestalt statt, d.h. es wird eine Gewichtsverlagerung innerhalb der Gestalt vorgenommen.

Neben dem für die Definition metrischer Texte konstitutiven Metrum sind bei der Produktion metrischer Texte häufig noch eine Reihe weiterer Regeln, z.B. Reim- oder Alliterationsregeln von Bedeutung. Alle diese Regeln können als Teil einer Grammatik metrischer Texte angesehen werden. Die Regeln dieser Grammatik operieren grundsätzlich auf der Basis von Sprache³⁰ und erzeugen Strukturen, die global als metrische Strukturen bezeichnet werden sollen.

Im folgenden soll nun das Verhältnis von metrischen und sprachlichen Strukturen genauer untersucht werden.

6. DIE INTERDEPENDENZ VON METRISCHEN UND SPRACHLICHEN STRUKTUREN

6.1. Das HEMPEL-OPPENHEIM-Schema einer wissenschaftlichen Erklärung

Die Untersuchung des Verhältnisses von metrischen und sprachlichen Strukturen ist deswegen für eine Theorie der Metrik von besonderer Bedeutung, weil eine solche Fragestellung die Möglichkeit eröffnet, nicht nur deskriptiv festzustellen, daß z.B. eine Sprache L ein metrisches System x mit den Eigenschaften $a, b, c...$ aufweist, sondern auch eine Antwort auf die Frage zu geben, *w a r u m* dies der Fall ist. Mit anderen Worten: Eine solche Fragestellung ermöglicht, theoretische Aussagen zu formulieren, die zur *E r k l ä r u n g* für das Vorliegen bestimmter Sachverhalte benutzt werden können. Die Untersuchung des Verhältnisses von metrischen und sprachlichen Strukturen ist somit ein wichtiger Schritt auf dem Weg zu einer empirischen Metriktheorie (vgl. Kap. 2.2).

6.1.1. Der Begriff der Erklärung ist einer der zentralen Begriffe der Analytischen Wissenschaftstheorie. Es herrscht jedoch keineswegs Einigkeit darüber, was unter einer wissenschaftlichen Erklärung¹ zu verstehen ist. Am häufigsten ist in der Analytischen Wissenschaftstheorie das auf HEMPEL/OPPENHEIM (1948) zurückgehende sog. HEMPEL-OPPENHEIM-Schema der Erklärung (H-O-Schema) diskutiert worden. Nicht selten wird sogar die Auffassung vertreten, daß alle wissenschaftlichen Erklärungen diesem Schema folgen sollten.

Im folgenden soll das H-O-Schema kurz charakterisiert werden. Ich folge dabei STEGMÜLLER (1969) - allerdings unter Verzicht auf eine Problematisierung.

Nach HEMPEL/OPPENHEIM hat eine wissenschaftliche Erklärung folgende allgemeine Struktur (vgl. STEGMÜLLER 1969:86).

	$A_1, \dots, A_n$
Explanans	
	$G_1, \dots, G_r$
Explanandum	<u>E</u>

$A_1, \dots, A_n$ sind Sätze, welche die Antezedensbedingungen beschreiben, $G_1, \dots, G_r$ sind allgemeine Gesetzmäßigkeiten. Diese beiden Klassen von Aussagen bilden das sog. Explanans. E ist die Beschreibung des zu erklärenden Ereignisses (Sachverhaltes), das sog. Explanandum. Es gelten folgende Adäquatheitsbedingungen:²

- B_1 "Das Argument, welches vom Explanans zum Explanandum führt muß k o r r e k t sein.
- B_2 Das Explanans muß m i n d e s t e n s e i n a l l g e m e i n e s G e s e t z enthalten. (oder einen Satz, aus dem ein allgemeines Gesetz logisch folgt).
- B_3 Das Explanans muß einen e m p i r i s c h e n Gehalt besitzen.
- B_4 Die Sätze, aus denen das Explanans besteht, müssen w a h r sein." (STEGMÜLLER 1969:86)

Die Bedingung B_1 bedeutet, daß bei deduktiv-nomologischen Erklärungen, das Explanandum l o g i s c h aus dem Explanans folgen muß. Bei induktiv-statistischen Erklärungen handelt es sich dagegen um einen Wahrscheinlichkeitsschluß, der ebenfalls bestimmten Adäquatheitsbedingungen genügen muß.

Soll ein Ereignis (Sachverhalt) nicht erklärt, sondern vorhergesagt werden, kann ebenfalls das H-O-Schema zugrundegelegt werden. Denn nach HEMPEL/OPPENHEIM ist jede Erklärung auch eine potentielle Prognose. Allerdings setzt diese Sicht die Strukturidentität von Erklärung und Prognose voraus, eine Voraussetzung, die keineswegs unproblematisch ist (vgl. STEGMÜLLER 1969:153ff).

6.1.2. Wir haben gesehen, daß bei einer adäquaten Erklärung das Explanans mindestens ein allgemeines Gesetz enthalten muß. Hier ergibt sich wiederum die Schwierigkeit, daß auch in bezug auf den Gesetzesbegriff in der Analytischen Wissenschaftstheorie keine Einigkeit herrscht (vgl. z.B. KRÜBER 1968; STEGMÜLLER 1969:273ff).

Als Definitionsmerkmale von Gesetzen werden u.a. der Raum-Zeit-Bezug und der Grad der Bewährung genannt. Beim augenblicklichen Stand der Theoriebildung in der Metrik erscheint es zweckmäßig, einen relativ schwachen Gesetzesbegriff zugrunde zu legen. Es soll deshalb jede gut bestätigte Hypothese, also auch eine gut bestätigte singuläre (raum-zeitlich gebundene) Hypothese - ALBERT (1957) nennt eine solche Hypothese ein Quasigesetz - im Explanans zugelassen sein.

Im folgenden sollen nun zuerst einige gesetzmäßige Beziehungen zwischen metrischen und sprachlichen Strukturen aufgezeigt werden. Anschließend soll anhand von einigen Beispielen dargelegt werden, wie das H-O-Schema in der Metrik angewendet werden kann.

6.2. Der Kontakt zwischen metrischen und sprachlichen Strukturen

Bei der Untersuchung der Interdependenz zwischen metrischen und sprachlichen Strukturen muß zwischen den Ebenen des Systems, der Norm und der Rede unterschieden werden.³ Metrische Strukturen bilden Systeme mit eigenen Realisationsnormen und sind, wie die Untersuchung der interlingualen Übertragung metrischer Systeme zeigt, zumindest partiell vom Sprachsystem unabhängig.⁴

Von der Untersuchung der Interdependenz von metrischen und sprachlichen Strukturen auf der Ebene des Systems ist die Untersuchung des direkten Kontakts zwischen metrischen und sprachlichen Strukturen auf der Ebene der Rede zu unterscheiden (vgl. HARWEG/SUERBAUM/BECKER 1969:399). Dieser Kontakt führt oft dazu, daß die sprachlichen Strukturen von den metrischen Strukturen affiziert werden.

Z.B. wird sehr häufig aufgrund des Metrums die normalsprachliche Wortstellung durchbrochen. Sogar die Struktur eines einzelnen Wortes kann, wie z.B. bei der im Griechischen vorkommenden Tmesis, zerstört werden. Das Metrum legt auch starke Beschränkungen bei der Wortwahl auf. So kann im lateinischen daktylischen Hexameter das für die römische Antike wichtige Wort *humanitas* aufgrund seiner

Silbenstruktur (Schema: --x-) nicht vorkommen. Im französischen Alexandriner werden, wie GUIRAUD (1953:30) statistisch nachgewiesen hat vor allem kurze Wörter bevorzugt, um einer Verseinheit (z.B. dem *hémistiche*) auch eine gedanklich-syntaktische Einheit entsprechen zu lassen.

Auch die Tempusverwendung im französischen Rolandslied ist vermutlich durch die Versstruktur bedingt. Dort findet sich das im Verhältnis zum *présent* und zum *passé simple* längere *passé composé* vor allem im längeren zweiten Versteil, während das *présent* und das *passé simple* vor allem im kürzeren ersten Versteil vorkommt (vgl. STEFENELLI-FÜRST 1966:80).

Häufig führt der Verszwang auch zur Verwendung von *nonsense*-Elementen, wie z.B. 'fallera' im Deutschen. Dies kann sogar, wie in manchen samojedischen Liedern, zur Folge haben, daß der Text kaum mehr verständlich ist (vgl. KOCH 1966:42f).⁵

6.3. Die Abhängigkeit des metrischen Systems vom Sprachsystem

Von größerem Interesse für eine Theorie der Metrik ist jedoch die Frage, inwieweit metrische Systeme vom Sprachsystem abhängen oder mit den Worten von IHWE (1975:379): "The way in which the structure of natural language permits the introduction and development of metrical systems and makes them thus actually possible." Für IHWE ist dies sogar das eigenliche Explanandum der generativen Metrik.

Es wird immer wieder betont, daß metrische Systeme auf den funktional relevanten Eigenschaften der Sprache basieren (vgl. z.B. STANKIEWICZ 1964; BECK 1966; LOTZ 1974). Dies ist schon 1933 von JAKOBSON deutlich formuliert worden:

"Nicht der Schall, sondern der Sprachlaut als solcher wird als Baustein des Verses verwendet. Das Relevante an den Sprachlauten ist ihr *phonologisch*er Wert, oder mit anderen Worten, diejenigen lautlichen Eigenschaften, welche in der gegebenen Sprache zur Sinnesunterscheidung dienen können. Erst dieser Wert macht die Laute zum Sprach- bzw. Versbestande." (JAKOBSON 1966:51)

Zur Erläuterung mögen einige wenige Beispiele genügen. Im Lateinischen und Griechischen ist die Quantität phonologisch relevant. Auch das Metrum basiert auf der Silbenquantität.⁶ Im Deutschen und Englischen ist, wie auch in vielen anderen Sprachen, der Akzent phonologisch relevant. Die Basis des Metrums ist ebenfalls der Akzent.⁷

In der Sprache als *langue* werden meist binäre Distinktionen getroffen⁸, beim Akzent z.B. zwischen betont - unbetont, bei der Quantität zwischen Länge und Kürze. Auf der Ebene der *parole* manifestieren sich diese Distinktionen dann in zahlreichen Abstufungen. Sprachen, die auf der Ebene des Systems mehrstufige Distinktionen benutzen, sind weitaus seltener. Hierzu zählen lediglich einige Tonsprachen (z.B. Chinesisch). Eine mehrstufig quantitative Opposition scheint es lediglich im Estnischen zu geben.⁹

Das Metrum scheint ausschließlich binäre Distinktionen zu verwenden. So wird im metrischen System des Chinesischen lediglich die Opposition zwischen gleichbleibenden und nicht gleichbleibenden Tönen gemacht (vgl. LOTZ 1964:140f). Die Binarität von Oppositionen scheint somit zumindest eine universale Tendenz, wenn nicht sogar ein metrisches Universale zu sein (vgl. LOTZ 1974:979).

Betrachtet man die interkulturelle Übertragung metrischer (Teil)systeme, wie z.B. des antiken Hexameters in das Deutsche oder Englische¹⁰, liegt es nahe, nicht einzelsprachliche prosodische Merkmale wie Länge oder Betonung als konstitutiv anzusehen, sondern die abstrakte Distinktion zwischen markiert und unmarkiert. So kann z.B. 'der' Hexameter auf einer abstrakten tiefenstrukturellen Ebene als Folge von markierten und unmarkierten Einheiten aufgefaßt werden, die einzelsprachlich durch die jeweiligen funktional relevanten Einheiten, wie z.B. lange bzw. kurze Silben im Lateinischen und Griechischen oder betonte bzw. unbetonte Silben im Deutschen oder Englischen, realisiert werden.

6.4. Zur diachronen Entwicklung metrischer Systeme

6.4.1. Da einzelsprachliche metrische Systeme auf den funktional relevanten Eigenschaften der jeweiligen Sprache basieren, führt eine Veränderung dieser Eigenschaften zumindest langfristig auch zu einer Veränderung des metrischen Systems. Die Richtung der Veränderung kann jedoch häufig nicht allein aufgrund der diachronen Entwicklung des Sprachsystems erklärt werden.

Metrische Systeme weisen, wie wir gesehen haben, einen hohen Grad an Konventionalität auf. Deshalb müssen in der Regel auch außersprachliche Faktoren, die häufig lediglich anhand einer komplexen interkulturellen Analyse ermittelt werden können, zur Erklärung herangezogen werden (vgl. IHWE 1975:380f).

6.4.2. Für die Abhängigkeit des metrischen Systems von der Entwicklung des Sprachsystems lassen sich zahlreiche Belege finden.¹¹

So ist, wie TARLINSKAJA (1974) zeigt, die Ersetzung der Alliteration durch den Reim in der zweiten Hälfte des 12. Jahrhunderts in England offensichtlich u.a. durch die Abschwächung des Wortakzents bedingt. Daneben spielt aber auch als außersprachlicher Faktor der Einfluß des romanischen Verses eine Rolle.

Ein illustratives Beispiel für die Abhängigkeit des metrischen Systems vom Sprachsystem ist die Geschichte des französischen Verses.

Als im Spätlateinischen aufgrund des vulgärlateinischen Quantitätenkollaps auch in der Dichtung das Gefühl für die Silbenquantität immer mehr verloren ging, hat sich das quantifizierende metrische System des Lateinischen zu einem isosyllabisch-akzentuierenden System entwickelt, das für die mittellateinische Dichtung konstitutiv wurde (vgl. CRUSIUS/RUBENBAUER 1963:130f). Unter dem Einfluß des mittellateinischen Systems hat sich dann im Französischen der Isosyllabismus als konstitutives, der Sprache angemessenes Prinzip herausgebildet (vgl. RAUHUT 1935).

In der weiteren Entwicklung des Französischen wurde jedoch der Isosyllabismus durch den Wandel des unbetonten *e* zum *a-instable*,

der zu einem Widerspruch zwischen normativer Metrik und Alltagssprache führte, in Frage gestellt. Dieser Widerspruch hatte zur Folge, daß der Isosyllabismus, wie bereits LOTE (1919) experimentell nachgewiesen hat, bei der lautlichen Aktualisierung von Versen häufig nicht eingehalten wird, weil das *à-instable* zur Hervorhebung bestimmter Wörter oft dort gelesen wird, wo es nach der normativen Metrik nicht erlaubt ist (vgl. auch MAROUZEAU 1955). LOTE (1919:701) stellt deshalb fest:

"Le syllabisme est un leurre, car la déclamation ne se soucie ni de conserver à l'alexandrin les douze ou treize éléments qui doivent le constituer, ni d'augmenter ou de diminuer ce nombre..."

Da der Isosyllabismus als konstitutives Prinzip oft als nicht ausreichend empfunden wird, müssen im Französischen in der Regel noch Reim und Zäsur als weitere metrische Konstituenten hinzukommen (vgl. GUIRAUD 1953:27).¹²

In der Renaissance - und hier zeigt sich wiederum der Einfluß außersprachlicher Faktoren - ist mehrmals der Versuch unternommen worden, entsprechend dem antiken Vorbild quantifizierende französische Verse zu schreiben. Auch DU BELLAY hält in seiner "Deffence et illustration de la langue francoyse" (vgl. Kap. I,9) den "vers prosodique" prinzipiell für möglich, weil er keinen Zusammenhang zwischen der Sprachstruktur und der Metrik sieht. Der quantifizierende Vers ist im Französischen jedoch deswegen zum Scheitern verurteilt, weil im Französischen im Gegensatz zum Lateinischen ein *rendement fonctionnel* der Quantität praktisch nicht existiert¹³ und die Quantität somit auch nicht als Funktionsträger im metrischen System benutzt werden kann, ohne daß der Vers als 'unfranzösisch' empfunden wird.

Aber auch eine rein akzentuierende Nachbildung antiker Verse war im Französischen im Gegensatz zum Deutschen oder auch zum Italienischen aufgrund des festen oxytonen Akzents und der Dominanz des Satzakkzents über den Wortakzent nicht möglich (vgl. ELWERT 1966:32f).

Auch im Deutschen versuchte man im Anschluß an das antike Vorbild, ohne den akzentuierenden Charakter der deutschen Sprache zu berücksichtigen, das quantitative Prinzip zu übertragen. Dies geschah ebenso wie in Frankreich vor allem in der Renaissance. Erst durch OPITZ "Buch von der deutschen Poeterey" (1624) wurde das akzentuierende Prinzip als der deutschen Sprache am angemessensten weitgehend anerkannt (vgl. HEUSLER 1956).¹⁴

6.5. Zur Anwendung des HEMPEL-OPPENHEIM-Schemas in der Metrik

Die bisherigen Bemerkungen erlauben es bereits, einige Gesetzmäßigkeiten zu formulieren, die für Erklärungen nach dem H-O-Schema benutzt werden können.

So scheint z.B. folgende allgemeine Hypothese gut bestätigt zu sein:

G_1 : Wenn in einer Sprache L ein prosodisches Merkmal P funktional relevant ist, ist P auch im metrischen System M_L dieser Sprache funktional relevant.

Um zu gewährleisten, daß das Argument, das vom Explanans zum Explanandum führt, korrekt ist (Adäquatheitsbedingung B_1), ist es zweckmäßig, die bei einer Erklärung verwendeten Aussagen zu formalisieren.

Nimmt man tentativ an, daß G_1 einer materialen Äquivalenz entspricht, kann G_1 folgendermaßen formalisiert werden:

$$G_1: P(L) \leftrightarrow P(M_L)$$

Mit Hilfe einer Definition wie

$$D_1: P := A \vee Q \vee \dots,$$

wobei A für 'Akzentuierung' und Q für 'Quantität' steht, lassen sich aus G_1 folgende speziellere Gesetze ableiten:

$$G_{11}: A(L) \leftrightarrow A(M_L)$$

$$G_{12}: Q(L) \leftrightarrow Q(M_L)$$

Diese Gesetze erlauben eine Reihe von Erklärungen. So kann z.B. gefragt werden, warum im metrischen System des Deutschen (D) die Akzentuierung funktional relevant ist. Dieses Explanandum kann, sofern man G_{11} als deterministisches Gesetz ansieht, mit Hilfe des *modus (ponendo) ponens* folgendermaßen abgeleitet werden:

$$G_{11}: \forall L \forall M_L [A(L) \leftrightarrow A(M_L)]$$

$$A_{11}: \frac{A(D)}{A(M_D)}$$

Bei dieser Erklärung handelt es sich um eine deduktiv-nomologische Erklärung. Nimmt man jedoch an, daß G_{11} als statistisches Gesetz anzusehen ist, das mit einer bestimmten Wahrscheinlichkeit p gilt, liegt eine induktiv-statistische Erklärung vor. In diesem Fall kann beim Vorliegen der Antezedensbedingung $A(D)$ aufgrund des Gesetzes G_{11} mit einer induktiven Wahrscheinlichkeit p auf das Vorliegen des Explanandums $A(M_D)$ geschlossen werden.

Ich werde im folgenden grundsätzlich davon ausgehen, daß es sich bei den entsprechenden Gesetzen um statistische Gesetze handelt. Hierfür sprechen vor allem zwei Gründe. Zum einen dürfte es in der Metrik ebenso wie in vielen anderen Bereichen kaum deterministische Gesetze geben. Zum anderen können deterministische Gesetze, falls sie auftreten sollten, unter den statistischen Gesetzesbegriff subsumiert werden, indem sie als statistische Gesetze mit der Wahrscheinlichkeit 1 bzw. 0 interpretiert werden.

Auch das Scheitern des quantifizierenden Verses im Französischen (F) oder im Deutschen hätte aufgrund der bisher formulierten Gesetze vorausgesagt werden können. Da nämlich die logische Äquivalenz

$$Q(L) \leftrightarrow Q(M_L) \leftrightarrow \neg Q(L) \leftrightarrow \neg Q(M_L)$$

und die Antezedensbedingung $\neg Q(F)$ gilt, gilt auch folgender Schluß:

$$G_{12}: \neg Q(L) \leftrightarrow \neg Q(M_L)$$

$$A_{12}: \neg Q(F)$$

$$\neg Q(M_F)$$

Auf die gleiche Weise kann auch die diachrone Veränderung metrischer Systeme erklärt werden. Schreiben wir für das zweistellige Prädikat 'verändert sich von x nach y' $V(x,y)$, kann z.B. mit Hilfe von G_1 folgendes allgemeine diachrone Gesetz formuliert werden:

$$G_2: V[P(L), \neg P(L)] \leftrightarrow V[P(M_L), \neg P(M_L)]$$

G_2 kann z.B. dazu benutzt werden, die Veränderung des metrischen Systems des Lateinischen von einem quantifizierenden zu einem nicht-quantifizierenden System zu erklären. Allerdings kann mit G_2 nicht die Richtung der Veränderung erklärt werden. Dazu bedarf es einer Reihe weiterer Gesetze.

Diese wenigen Beispiele mögen an dieser Stelle genügen. Ein weiterer Hinweis zur Anwendung des H-O-Schemas in der Metrik findet sich in Kapitel 13.

7. DIE SILBE ALS METRISCHE KONSTITUENTE

7.1. Wir haben gesehen, daß der Isosyllabismus in Verbindung mit Reim und Zäsur als konstitutiv für das metrische System des Französischen angesehen werden kann.

Neben Reim, Zäsur und Isosyllabismus weisen französische Verstexte jedoch noch eine Vielzahl weiterer spezifischer Charakteristiken auf. So finden sich z.B. Erscheinungen wie Alliteration oder syntaktische und semantische Parallelismen. Diese sind jedoch im Französischen nicht konstitutiv.¹

Dies gilt aber nicht für andere metrische Systeme. So kann die Alliteration (Stabreim) als strukturelle Konstituente der altgermanischen Langzeile angesehen werden (vgl. z.B. KURYŁOWICZ 1966, 1970) und der syntaktisch-semantische Parallelismus als Konstituente des metrischen Systems des Althebräischen (vgl. YODER 1972) oder des Ugaritischen (vgl. GORDON 1965).²

7.2. Von den metrischen Konstituenten ist die Silbe von besonderer Bedeutung. LOTZ (1974:972) meint sogar: "All strictly regulated metric systems are founded on syllabification."

Trotz der offensichtlichen Bedeutung der Silbe für die Metrik wird jedoch nur in wenigen Arbeiten zur Metrik der Versuch unternommen, das jeweils zugrundegelegte Verständnis des Terminus 'Silbe' zu explizieren. Anscheinend wird vorausgesetzt, daß jeder weiß, was unter diesem Terminus zu verstehen ist.³

In der Linguistik herrscht jedoch keineswegs Einigkeit in bezug auf die Silbe.⁴ Denn sowohl die Definition als auch der Stellenwert der Silbe in der linguistischen Theoriebildung⁵ ist umstritten. Konsens scheint lediglich darüber zu bestehen, daß "die Silbe eine ... in der Rede h ö r b a r e Einheit ist." (PILCH 1968:17).

Es sind bisher zahlreiche Versuche unternommen worden, dieser intuitiven Silbenvorstellung motorische, akustische, artikulatorische und perzeptorische Korrelate zuzuordnen. So ist die Silbe z.B. von SIEVERS (1893:183) und STETSON (1928) als Drucksilbe, von DE SAUSSURE (1968:86ff) als Öffnungssilbe oder von FOUCHÉ (1930:33) und von DURAND (1954) als Spannungssilbe definiert worden. Außerdem wird seit einiger Zeit auch nach einem psychischen Korrelat⁶ gesucht.

Das akustische und perzeptorische Korrelat wird von JAKOBSON/HALLE (1970:423) folgendermaßen beschrieben:

"In its acoustic aspect the crest usually exceeds the slopes in intensity and in many instances shows an increased fundamental frequency. Perceptually, the crest is distinguished from the slopes by a greater loudness, which is often accompanied by a heightened voice pitch."⁷

7.3. Keine der bis heute vorgeschlagenen phonetischen Silbendefinitionen erlaubt jedoch eine in jedem Fall eindeutige Festlegung der Silbengrenzen. Denn "trotz all der vielen Überlegungen und Versuche ist es den Phonetikern immer noch nicht gelungen, ein restlos zuverlässiges Kriterium silbischer Gliederung im Sprechvorgang zu finden." (VON ESSEN 1966:132, vgl. auch LADEFOGED 1975:222).

Aus dieser Tatsache ist die Konsequenz gezogen worden, die Silbe nicht als phonetische, sondern als phonologische Einheit anzusehen, deren Struktur durch die jeweiligen einzelsprachlichen Kombinationsregeln bestimmt ist (vgl. VON ESSEN 1951:200).⁸

Auf der Grundlage dieser Konzeption sind verschiedene Versuche gemacht worden, die intuitive Vorstellung von Silbengrenzen durch objektive Delimitierungsmethoden zu ersetzen.⁹ Ein umfassender Versuch zur Aufstellung nicht einzelsprachlich gebundener Regeln zur Silbentrennung ist von PULGRAM (1970) unternommen worden.

Aber erst LEHFELDT (1971) ist es durch die Einbeziehung statistischer Kriterien gelungen, auf der Basis der von PULGRAM vorgeschlagenen Regeln einen Algorithmus zu entwickeln, der eine von der Intuition des Untersuchenden unabhängige, automatische Silbendelimitation ermöglicht.

7.4. Das Problem der exakten, intersubjektiv überprüfbaren Bestimmung der Silbengrenzen, das z.B. bei typologischen Untersuchungen zur Silbenstruktur wegen des typologischen Isomorphismus von großer

Bedeutung ist, spielt in der Metrik lediglich eine marginale Rolle. Denn in der Regel ist hier nur die prosodische Beschaffenheit des Silbennucleus oder im Falle von isosyllabischen Systemen die Anzahl der Nuclei von Interesse.¹⁰ Der Silbennucleus ist jedoch, wie wir gesehen haben, sowohl akustisch als auch perzeptorisch relativ leicht zu bestimmen.

Für die Metrik scheint deshalb folgende, von LOTZ (1964:138) vorgeschlagene Silbendefinition praktikabel zu sein: "It is possible to use syllable as a position of sound determined by one syllabic crest, where the boundaries are not necessarily definite."¹¹

Für manche Fragestellungen in der Metrik ist noch die terminologische Unterscheidung zwischen der 'emischen' Silbe, dem Syllabem, und der 'etischen' Silbe, dem Syllab, von Interesse (vgl. PIKE 1967: 365; HAMMARSTRÖM 1966).

HAMMARSTRÖM (1966:38f) definiert das Syllab und das Syllabem folgendermaßen:

"Wir schlagen vor, jedes Beispiel einer ausgesprochenen Silbe ein S y l l a b zu nennen (...) Syllabe werden in Klassen, S y l l a b e m e n, zusammengefaßt: Syllabe, die an der gleichen Stelle in verschiedenen Beispielen des gleichen Wortes vorzufinden sind, gehören dem gleichen Syllabem an. (...) Es wird angenommen, daß sich Syllabe zu Syllabemen wie Phone zu Phonemen verhalten. (...) Das Syllabem besitzt bestimmte S e g m e n t e, worunter Phonemrepräsentanten, die den Phonem des Syllabs entsprechen, verstanden werden."

Die Unterscheidung zwischen Syllabem und Syllab ist u.a. deshalb wichtig, weil nur das Syllabem als Folge von Phonemrepräsentanten metrisch konstitutiv sein kann.¹²

8. DER REIM

Der Reim findet sich zwar auch in Prosatexten, so z.B. als sog. Homoioteleuton in der spätantiken Kunstprosa, er wird jedoch vor allem in metrischen Texten verwendet. Dort ist er eines der auffallendsten Charakteristiken der Ausdrucksseite des Textes.

Der Reim ist einmal ein sprachliches Phänomen, das z.B. mit den Methoden der Phonologie untersucht werden kann, zum andern eine literarisch-ästhetische Konvention. Er hängt somit sowohl vom jeweiligen Sprachsystem als auch von der jeweiligen literarisch-ästhetischen Tradition ab.

In der Phonologie ist der Reim bisher vor allem im Zusammenhang mit der Rekonstruktion des phonologischen Systems älterer Sprachstufen untersucht worden (vgl. z.B. MARTHE 1968). Synchrone phonologische Untersuchungen finden sich dagegen relativ selten (vgl. z.B. KRÄMER 1971, 1974). Auch die Frage nach dem Stellenwert des Reims in der Theoriebildung der Phonologie ist bisher erst vereinzelt erörtert worden (vgl. z.B. LUELSDOORFF 1968).

8.1. Zur Definition des Reims

8.1.1. Unter einem Reim soll im folgenden eine binäre Relation R auf dem Vokabular VO einer Sprache L verstanden werden. R ist somit eine Teilmenge des kartesischen Produkts $VO \times VO$. Die Elemente von R sind jeweils geordnete Paare von (partiellen) Homophonen, bzw. im Fall des Augenreims von (partiellen) Homographen aus dem Vokabular VO .

Wir wollen die Relation R jetzt als eine Äquivalenzrelation betrachten und untersuchen, inwieweit diese Konzeption der jeweiligen einzelsprachlichen literarischen Tradition entspricht (vgl. ABERNATHY 1967).

Als Äquivalenzrelation muß R reflexiv, symmetrisch und transitiv sein, d.h. es muß für alle $x, y, z \in VO$ gelten:

1) Reflexivität

$$\forall x [(x,x) \in R]$$

Jedes Wort reimt mit sich selbst.

2) Symmetrie

$$\forall x \forall y [(x,y) \in R \wedge (y,x) \in R]$$

Reimt ein Wort x mit einem Wort y, so reimt auch das Wort y mit dem Wort x.

3) Transitivität

$$\forall x \forall y \forall z [(x,y) \in R \wedge (y,z) \in R \Rightarrow (x,z) \in R]$$

Wörter, die mit demselben Wort reimen, reimen auch miteinander.

Der Reim als Äquivalenzrelation gilt jedoch nur sehr bedingt.¹ Häufig gilt z.B. die Eigenschaft der Reflexivität nicht. So ist z.B. im Französischen der identische Reim (*la rime du même au même*), der des öfteren im Altfranzösischen zu finden ist, von der Plejade an verpönt. Erst seit dem Symbolismus wird er wieder verwendet (vgl. ELWERT 1966:94).

Die Relation R ist also im Französischen nicht totalreflexiv und zu bestimmten Epochen sogar irreflexiv.

Auch die Transitivität gilt nicht uneingeschränkt. So war z.B. im Französischen der Reim zwischen zwei Komposita mit gleichem Stamm (*rime dérivative*) zeitweise nicht erlaubt (vgl. ELWERT 1966:92ff). Eine ähnliche Konvention findet sich im Ungarischen (vgl. LOTZ 1942:131). Dadurch gilt die Transitivität zwar z.B. für die Reime

$$\text{malheur} \tilde{R} \text{ coeur} \wedge \text{coeur} \tilde{R} \text{ soeur} \Rightarrow \text{malheur} \tilde{R} \text{ soeur}$$

nicht jedoch für

$$\text{malheur} \tilde{R} \text{ coeur} \wedge \text{coeur} \tilde{R} \text{ bonheur} \not\Rightarrow \text{malheur} \tilde{R} \text{ bonheur.}$$

Die Relation R ist also im Französischen zu bestimmten Epochen lediglich teiltransitiv, d.h. die Transitivität kann gelten oder nicht.

Erfüllt eine Relation zwar die Bedingung der Reflexivität und der Symmetrie, nicht jedoch die Bedingung der Transitivität, so heißt eine solche Relation auch Ähnlichkeitsrelation oder Toleranzrelation (vgl. FISCHER 1973:25).

Beim sog. unreinen Reim ist weit häufiger als beim reinen Reim die Bedingung der Transitivität nicht erfüllt. So reimt zwar nach KRÄMER (1974:331) z.B. [ε:] mit [e:] und [e:] mit [ø:], nicht aber [ε:] mit [ø:]. Ähnliches gilt für den unreinen Reim in anderen Sprachen (vgl. ABERNATHY 1967). Da der unreine Reim somit grundsätzlich nur teiltransitiv zu sein scheint, ist beim unreinen Reim von vornherein davon auszugehen, daß keine Äquivalenzrelation, sondern höchstens eine Ähnlichkeitsrelation vorliegt.

Für bestimmte Fragestellungen ist es von Interesse, nicht nur die Relation R, sondern auch die komplementäre Relation $\bar{R}$ zu betrachten.²

So sind z.B. im chinesischen "regulated verse", der Hauptform des chinesischen Verses seit der T'ang Dynastie, sowohl die Reimpositionen festgelegt als auch die Positionen, die nicht reimen dürfen, d.h. es ist sowohl die Zugehörigkeit zu R als auch zum Komplement von R konventionalisiert (vgl. LIU 1966:26ff; ABERNATHY 1967; JAKOBSON 1970).

8.1.2. Die Äquivalenz von Reimwörtern beruht auf einer Identitäts- bzw. auf einer Ähnlichkeitsrelation zwischen jeweils zwei Lauten oder Lautfolgen, von denen mindestens ein Laut ein Vokal sein muß. Die Bedingung der Identität gilt für den sog. reinen Reim. Die schwächere Bedingung der Ähnlichkeit berücksichtigt auch den sog. unreinen Reim. Ist die Identitäts- oder Ähnlichkeitsrelation auf die 'Reimvokale' beschränkt, wird anstelle von Reim häufig auch der Terminus Assonanz verwendet. Soll explizit gemacht werden, daß es sich bei einem Reim nicht um eine Assonanz handelt, wird auch der Terminus Vollreim gebraucht.

Je nachdem, welche Position die identischen oder ähnlichen Laute oder Lautfolgen im Wort einnehmen, kann zwischen Endreim (ER), Anfangsreim (AR) und Binnenreim (BR) unterschieden werden.³

Bezeichnen wir die Laute oder Lautfolgen mit den Buchstaben $a, b, c, \dots$, die nicht-kommutative Relation der Verkettung als $\cap$ und die Ähnlichkeitsrelation als $\approx$, können diese Reimtypen folgendermaßen definiert werden:

$$1) \text{ ER}(x, y) :\Leftrightarrow \forall x \exists y [(x = a \cap b) \wedge (y = c \cap d) \wedge (b \approx d) \wedge \neg(a \approx c)]$$

Eine Endreimrelation besteht zwischen jedem Element x und einem Element y genau dann, wenn x und y Ketten von Lauten oder Lautfolgen sind, wobei die zweite Komponente jeder Kette jeweils ähnlich, die erste Komponente jeweils nicht ähnlich ist.

Beispiel: franz. *frémir* $\overset{\sim}{\text{ER}}$ *sentir*

Für die weit selteneren Fälle des Anfangs- und des Binnenreims gilt analog:

$$2) \text{ AR}(x, y) :\Leftrightarrow \forall x \exists y [(x = a \cap b) \wedge (y = c \cap d) \wedge (a \approx c) \wedge \neg(b \approx d)]$$

Beispiel: engl. *because* $\overset{\sim}{\text{AR}}$ *before*

$$3) \text{ BR}(x, y) :\Leftrightarrow \forall x \exists y [(x = a \cap b \cap c) \wedge (y = d \cap e \cap f) \wedge (b \approx e) \wedge \neg(a \approx d) \wedge \neg(c \approx f)]$$

Beispiel: rot $\overset{\sim}{\text{BR}}$ Mohn

Ersetzt man in den Definitionen 1 - 3 $\approx$ durch $=$, d.h. verschärft man die Ähnlichkeitsrelation zu einer Identitätsrelation, werden die Definitionen auf den reinen Reim eingeschränkt.

Interpretiert man a, b, c, d jeweils als Grapheme bzw. Graphemfolgen, werden durch die Definitionen 1 - 3 verschiedene Formen des Augenreims definiert.

8.1.3. In einem Text müssen die reimenden Einheiten in einer Rekurrenzrelation (Rek) stehen. Benutzen wir das Symbol $<$ für die Vorgängerrelation und das Symbol $\sim$ für die Äquivalenzrelation, kann diese Relation folgendermaßen definiert werden:

$$\text{Rek}(x, y) :\Leftrightarrow \forall x \exists y (y < x \wedge x \sim y)$$

Der Reim ist somit ein Beispiel für das bekannte JAKOBSONsche Äquivalenzprinzip: "The poetic function projects the principle of equivalence from the axis of selection into the axis of combination." (JAKOBSON 1964:358).

In der Regel sind die Reimpositionen und, wie wir gesehen haben, zuweilen auch die Positionen, in denen nicht gereimt werden darf, aufgrund der literarischen Tradition fixiert.

Gerade in der Volkspoesie findet sich jedoch auch häufig eine Tendenz zum Reim in 'Nicht-Reimpositionen'. So besteht z.B. im malaischen Pantun eine statistisch signifikante Tendenz zur Assonanz in parallelen Positionen innerhalb eines Verses (vgl. ALTMANN 1963).

Aufgrund der Reimpositionen können noch zwei weitere Reimtypen unterschieden werden. Befinden sich die Reimwörter in verschiedenen Versen (V), d.h. gilt für die Komponenten der geordneten Paare (x, y) die Bedingung $x \in V_i, y \in V_j$ und $i \neq j$, soll der Reim als *vertikaler Reim* bezeichnet werden. Eine solche vertikale Reimrelation besteht meist zwischen Wörtern am Versende.

Weit seltener ist der *horizontale Reim*, bei dem die Bedingung $x, y \in V_i$ gilt. Ein Beispiel hierfür ist die sog. *rime couronnée* in den folgenden Versen MOLINETs:

Guerre a fait maint *chatelet laid*
Et mainte bonne *ville vile*

(zit. ELWERT 1966:114)

Allerdings werden solche Reime von der französischen Literaturkritik schon seit langem als Spielerei abgelehnt.

In bezug auf die Reimposition gilt noch eine wichtige Einschränkung.

Zusammengehörende Reime werden in der Sprechdichtung meist durch höchstens drei dazwischengeschobene Verszeilen voneinander getrennt, da sonst das Reimschema nicht mehr deutlich perzipiert werden kann. In gesungener Dichtung kann der Abstand dagegen größer sein, da es möglich ist, die Reimbeziehung durch die Melodie zu verdeutlichen (vgl. BAEHR 1970:78).

8.2. Zur Semantik des Reims

8.2.1. Die bisherigen Bemerkungen zum Reim bezogen sich lediglich auf die Lautebene. Der Reim kann aber auch auf der morphologischen, syntaktischen und semantischen Ebene analysiert werden, da zwischen den Reimelementen auch auf diesen Ebenen Identitäts- oder Ähnlichkeitsrelationen bestehen können.⁴ Auf der semantischen Ebene handelt es sich dabei meist nicht um eine denotative, sondern um eine konnotative⁵ Beziehung.

Eine solche konnotative Beziehung⁶ findet sich nach COHEN (1966:221) zwischen den Wörtern *soeur* und *douceur* in den folgenden BAUDELAIRE-Versen

Mon enfant, ma soeur
songe à la douceur

Welche Art von Konnotationen jeweils mit den Reimwörtern verbunden sind, hängt von einer Reihe von Faktoren ab. So verbindet z.B. GARY-PRIEUR (1971:101) mit *soeur* "connotations amoureuses, incestueuses", während für A. MARTINET diese Konnotationen aufgrund seiner Erziehung nicht akzeptabel sein sollen (vgl. GARY-PRIEUR 1971:101).

Legt man einen weiten Begriff von Konnotation zugrunde, der auch den Begriff der Assoziation mit einschließt (vgl. KLOEPFER 1975:91), dürften die jeweiligen Konnotationen auch z.B. durch geschlechts-, bildungs- oder entwicklungsspezifische Differenzen im Assoziationsverhalten der Perzipienten bedingt sein (vgl. HÖRMANN 1970:115 ff, LIST 1972:36 ff).

Da außerdem die Habit-Stärke von Assoziationen auch vom Kontext der Stimuluswörter abhängt, ist auch ein Einfluß des Kontexts auf die Reimkonnotationen zu vermuten.

Häufig werden durch den Reim auch Wörter verbunden, die für die meisten Rezipienten keine konnotativen Beziehungen aufweisen. In diesem Fall werden erst durch den Autor mögliche neue Konnotationen suggeriert.

In bezug auf die Wirkung des Reims kann vermutet werden, daß gerade solche Reime, die semantisch und/oder syntaktisch ungleiche Einheiten verbinden, einen hohen Überraschungswert haben (vgl. OOMEN 1973:67). Außerdem dürfte der Überraschungswert auch noch von der allgemeinen poetischen Kompetenz und vom "Erwartungshorizont" (JAUSS 1959) der Rezipienten abhängen.

Diese Vermutungen könnten mit Hilfe eines psychologischen Tests z.B. nach der Methode der Cloze-Procedure experimentell überprüft werden (vgl. z.B. TAYLOR 1953; FAULSTICH 1976).

8.2.2. Durch den Reim kann auch die semantische Dekodierung eines Textes erschwert werden (vgl. COHEN 1966:78 ff).

Dies trifft vor allem auf die "rime totale" zu - COHEN (1965:93) bezeichnet so homophone Verse vom Typ

Gal, amant de la Reine, alla, tour magnanime,
Galamment de l'arène à la tour Magne à Nîmes.

- und auch auf die vor allem bei den französischen 'Rhétoriciens' häufigen *rimes équivoques* vom Typ

Drap blanc, satin cardinalice,
Dans l'ombre du car d'ine Alice.
(J. PELLERIN)⁷

Die *rime totale* ist aber äußerst selten und kann wohl als Spielerei angesehen werden. Auch die *rime équivoque* ist seit der Klassik nur noch dann erlaubt, wenn sie aus einem einzigen Wort besteht und wenn die beiden reimenden Wörter deutlich als verschieden empfunden werden (ELWERT 1966:95). Deshalb ist ein Reim wie

le soir tombe: vers la tombe

(V. HUGO, Légende des siècles)⁸

erlaubt, da die Disambiguierung durch die Verschiedenheit der Kontextklassen erleichtert wird. So bezeichnet auch MEYER-EPPLER (1969:437) nur solche gleichlautenden Wörter als "äquivok" oder "homonym", die der gleichen Kontextklasse angehören.

Bei den übrigen Formen des Reims wird zwar ebenfalls durch die geringe phonemische Distanz der Reimwörter die Störbarkeit des Kommunikationsaktes erhöht,⁹ es dürfte jedoch selten zu tatsächlichen Störungen bei der Dekodierung kommen.

8.3. Reiner und unreiner Reim

Ich habe bei der Unterscheidung zwischen reinem und unreinem Reim darauf verzichtet, die Bedingungen anzugeben, aufgrund deren Laute bzw. Lautfolgen als identisch oder als ähnlich klassifiziert werden können. Die meisten dieser Bedingungen sind jedoch nicht universal, sondern hängen von der jeweiligen einzelsprachlichen literarischen Tradition ab. Ich werde mich deshalb auf einige grundsätzliche Feststellungen beschränken.

8.3.1. Für den reinen Reim nehme ich phonologische Identität, für den unreinen Reim lediglich phonetische Ähnlichkeit an.¹⁰ Phonetische Ähnlichkeit bedeutet dabei nicht Ähnlichkeit zwischen Phonemen, d.h. zwischen lautlichen Realisationen, sondern Ähnlichkeit zwischen Phonemen aufgrund einer partiellen Übereinstimmung ihrer Merkmale.

Phonetische Differenzen bei phonologischer Identität sind z.B. bei besonderen Sprechsituationen auf der Bühne denkbar. Ein Reimpaar kann z.B. bei dialogisch-stichomythischem Sprechen durch einen männlichen und einen weiblichen Sprecher realisiert werden. Außerdem kann ein Sprecher laut, der andere leise sprechen. Oft kann ein Sprecher auch stilistische Varianten eines Phonems benutzen, so daß eine allophonische Differenz bei phonologischer Gleichheit entsteht (vgl. KRÄMER 1971:27). Abgesehen von dieser "interindividuellen Variation" wird der gleiche Sprecher aufgrund der "intra-individuellen Variation" niemals bei der Aussprache von Reimen voll-

ständige phonetische Identität erzielen können (vgl. HEIKE 1972:23).

Von dieser Art der Variation, die zwar bei einer empirischen Untersuchung der Wirkung von Reimen berücksichtigt werden muß, ist bei der Definition des Reims als textuelles Phänomen jedoch zu abstrahieren.

Die Forderung nach Identität gilt aber nicht nur auf der Ebene des Phonems, sondern auch auf suprasegmentaler Ebene.

Im Italienischen, Spanischen und Deutschen z.B., die zu den Sprachen zählen, in denen die Akzentstelle phonologisch relevant ist,¹¹ gilt für den reinen Reim noch die zusätzliche Forderung, daß die identische Lautfolge mit dem letzten hauptbetonten Vokal beginnen muß (vgl. KRÄMER 1971:13; ELWERT 1968:86; BAEHR 1962:32). Deshalb gilt der Reim in den folgenden Versen von RINGELNATZ auch nicht als rein.¹²

Und jeder Beifall stärkte meinen Schwung
Die Dicke schwieg. Ich gab die Vorstellung.
(zit. KRÄMER 1971:17)

Dieses Phänomen hat STETSON schon 1903 untersucht und festgestellt: "Rhymes of syllables which have accents of strikingly different degrees are difficult to feel." (STETSON 1903:430).

So werden z.B. Reime wie

_____ dó	oder	_____ dō	oder	_____ gō
_____ gō		_____ gó		_____ dō
				_____ dō

nur mit Schwierigkeit als Reime perzipiert. Allerdings sind diese Ergebnisse möglicherweise nur beschränkt valide, da STETSON seine Untersuchungen lediglich an sinnlosen Silben vorgenommen hat.

8.3.2. Sowohl phonologische Identität als auch phonetische Ähnlichkeit sind lediglich notwendige, jedoch keine hinreichenden Bedingungen. Im Französischen z.B. wurde bis zum Symbolismus für den reinen Reim auch eine Berücksichtigung des Schriftbildes gefordert (vgl. ELWERT 1966:132 f). So konnte zwar *coup* mit *loup*,

jedoch nicht mit *tout* reimen, obwohl in allen drei Wörtern die Auslautkonsonanten schon seit dem Mittelalter nicht mehr gesprochen werden.

Das Italienische bietet sogar ein Beispiel dafür, daß ein nach der vorgeschlagenen Definition unreiner Reim als rein empfunden wird. So gelten bei den italienischen Verstheoretikern Reime auf /o/ : /o/ und /e/ : /e/ als rein (vgl. ELWERT 1968:80), obwohl der Unterschied im Öffnungsgrad im Italienischen in betonter Silbe phonologisch relevant ist.

Der Grund für diese Erscheinung ist in der literarischen Tradition zu suchen. Die sizilianische Dichterschule war das literarische Vorbild der Toskaner. Im fünfstufigen sizilianischen Lautsystem gab es aber nur die Phoneme /o/ und /e/ (vgl. LAUSBERG 1963:149). So konnten die Toskaner einen sizilianischen Reim wie *egre* : *amore* als *egre* : *amore* übernehmen, im Glauben, daß im Sizilianischen /o/ und /o/ reimten.¹³

8.4. Der unreine Reim

8.4.1. Der unreine Reim ist im Gegensatz zum reinen Reim weit weniger konventionalisiert. In den Poetiken oder Verslehren finden sich entweder überhaupt keine oder wenig eindeutige Vorschriften.

Außerdem ist der unreine Reim weit seltener als der reine Reim. KRÄMER (1971) fand bei einer Stichprobe von 58718 Versen von acht deutschsprachigen Autoren (BUSCH, RINGELNATZ, RILKE, MÖRIKE, HÖLDERLIN, MORGENSTERN, NIETZSCHE, TRAKL) nur durchschnittlich 13,64 % unreine Reime.¹⁴

Bei den Vokalen hat KRÄMER (1971:24f) sowohl quantitative als auch qualitative Differenzen festgestellt. Bei den qualitativen Differenzen überwiegen die Unterschiede im Öffnungsgrad, wobei die Öffnung selten über einen Grad hinaus verändert wird. Differenzen in der Artikulationsstelle sind seltener. Artikulatorische Sprünge, wie von *front* nach *back*, kommen zwar vor, sind aber aufgrund ihrer niedrigen Häufigkeit nicht repräsentativ. Bei den Konsonanten finden sich weit häufiger Differenzen in der Sonorität und im Verschluß als Differenzen in der Artikulationsstelle.

KRÄMER (1974) interpretiert diese Ergebnisse dahingehend, daß beim unreinen Reim nicht von einer Übereinstimmung zwischen Reimsegmenten auszugehen sei, sondern zwischen abstrakten Einheiten, die er poetische Grundphoneme oder Poetophoneme nennt. Differierende Reimsegmente wären dann jeweils als Realisationen ein und desselben Poetophonems anzusehen (vgl. auch SCHÄDLICH 1969).

8.4.2. Es ist mehrfach versucht worden, die Theorie der distinktiven Merkmale zur Erklärung von Differenzen zwischen Reimsegmenten heranzuziehen.

CAMARA (1953:126) meint, daß eine vollständige Identität der distinktiven Merkmale¹⁵ nicht nötig sei, so lange der Unterschied nur in "traços distintivos estrutural e funcionalmente fracos" bestehe.¹⁶ Es stellt sich die Frage, ob sich diese funktional und struktural schwachen Merkmale universal bestimmen lassen oder ob sie einzelsprachlich bestimmt werden müssen.

Eine Möglichkeit der Hierarchisierung der distinktiven Merkmale bietet sich aufgrund der Reihenfolge des Erwerbs bei der Spracherlernung des Kindes und des umgekehrt verlaufenden Verlustes im Falle der Aphasie (vgl. JAKOBSON/HALLE 1970:427). Auch die Universalität der Merkmale kann als Kriterium zur Hierarchisierung benutzt werden (vgl. JAKOBSON/HALLE 1970:429). Eine anerkannte Hierarchisierung aufgrund dieser oder weiterer Kriterien existiert bisher aber noch nicht. So sagen z.B. CHOMSKY/HALLE (1968:300):

"This subdivision of features is made primarily for purposes of exposition and has little theoretical basis at present. It seems likely, however, that ultimately the features themselves will be seen to be organized in a hierarchical structure which may resemble the structure that we have imposed on them for purely expository reasons."

Aufschlußreich ist auch die Betrachtung des unreinen Reims im slawischen Sprachkreis. Im Tschechischen z.B. können in der mündlichen Dichtung und in einigen Phasen der schriftlichen Dichtung Wörter wie *boty*, *boky*, *stopy*, *kosy*, *sochy* miteinander reimen,

während entgegengesetzt zum Deutschen der Reim zwischen stimmhaften und stimmlosen Konsonanten (*body* : *boty*) nicht erlaubt ist (vgl. JAKOBSON 1964:374). Ähnliches gilt auch für andere slawische Sprachen.

Diese Beispiele und auch die Feststellungen CAMARAs für das Brasilianische weisen darauf hin, daß der unreine Reim in Abhängigkeit von den "proprias condições fonémicas da língua" (CAMARA 1953:126) gesehen werden muß. Außerdem muß auch noch ein Einfluß der jeweiligen literarischen Konventionen auf die "Hierarchisierung der phonologischen Distinktionen" angenommen werden (vgl. JAKOBSON 1964:374).

Ein weiteres Kriterium für die Notwendigkeit einer einzelsprachlichen Bestimmung des unreinen Reims ist die Abhängigkeit der Perzeption von Phonemen vom Phonemsystem der entsprechenden Sprache. Denn der artikulatorische und der akustische Abstand zwischen zwei Lauten wird oft erst aufgrund der Untergliederung des Lautkontinuums in digitale Signale (Phoneme) perzipiert. So wird z.B. im Französischen der Abstand von /e/ zu /ɛ/ als weit geringer empfunden als der akustisch und artikulatorisch ungefähr gleiche Abstand zwischen /i/ und /e/, weil der Unterschied von /e/ und /ɛ/ - von manchen Sprechern wird überhaupt kein Unterschied mehr gemacht - neutralisierbar ist (FAURE 1964:326).¹⁷

8.4.3. Bei der Analyse des unreinen Reims sollte auch der auditive Eindruck mitberücksichtigt werden, da dieser möglicherweise mit darüber entscheidet, welche Merkmale als funktional schwach angesehen werden.

Hier ergibt sich jedoch das Problem, daß bei den üblicherweise verwendeten distinktiven Merkmalen der perzeptorische Aspekt weitgehend unberücksichtigt bleibt, obwohl die Namen einiger Merkmale, wie z.B. *acute* oder *grave* bei JAKOBSON/HALLE (1970), auf auditive Korrelate hinzuweisen scheinen. JAKOBSON/HALLE (1970:434) sagen jedoch selbst, daß nur der akustische und der motorische Aspekt der Merkmale hinreichend erforscht sei. Solange eine eindeutige intentionale¹⁸ Bestimmung von Eindrucksqualitäten noch nicht gelungen ist (vgl. HEIKE/THÖRMANN 1973:103), sagt somit eine Merkmalsanalyse nur sehr wenig über den auditiven Eindruck von Reimen aus.

Der auditive Eindruck ist möglicherweise auch der Grund dafür, daß bei dem von KRÄMER analysierten Korpus der Anteil der vokalischen Differenzen 83,66 % beträgt, obwohl die Vokalphoneme nur 44 % des Korpus ausmachen (vgl. KRÄMER 1971:23). Die betonten Reimvokale sind vermutlich aufgrund ihres im Verhältnis zu den Konsonanten weit größeren zeitlichen Anteils am Sprechereignis von entscheidender Bedeutung für den auditiven Eindruck des Reims. Wenn man von der Annahme ausgeht, daß die von KRÄMER analysierten Autoren den unreinen Reim als bewußtes Stilmittel eingesetzt haben, könnte als Grund für die Bevorzugung von Vokaldifferenzen die im Verhältnis zu den konsonantischen Differenzen größere auditive Wirkung angenommen werden.¹

Bei der Perzeption von Reimen spielen neben den lautlichen Eigenschaften der Reime auch noch eine Reihe weiterer Faktoren, wie z.B. die Bedeutung der Reimwörter oder die Reimerwartung, eine Rolle. Für den Einfluß der Reimerwartung läßt sich folgendes Beispiel anführen.

Die HÖLDERLIN-Verse

Wann der frohe Puls so plötzlich stünde
und verworren Freudesstimme tönte

(zit. KRÄMER 1971:19)

werden, aus ihrem Kontext gelöst, aufgrund der großen phonemischen Distanz der Reimwörter nicht als reimend empfunden.²⁰ Erst die durch die vorangehenden Reime geschaffene Erwartung läßt uns die Verse als gereimt empfinden.

8.5. Reim und Sprachstruktur

Zwischen Reim und Sprachstruktur bestehen eine Reihe von gesetzmäßigen Zusammenhängen, die ebenso wie die im Kap. 6 aufgezeigten Zusammenhänge Erklärungen ermöglichen.

Die Möglichkeit der Bildung von Reimen hängt zumindest in dreifacher Hinsicht von der Sprachstruktur ab. Einmal entstehen Beschränkungen aufgrund der phonologischen Struktur einer Sprache, zum anderen aufgrund der morphologischen Struktur. Außerdem bestimmt

auch der Umfang des Lexikons die Anzahl der Reimmöglichkeiten.

8.5.1. Das Portugiesische z.B. hat aufgrund seiner äußerst großen Anzahl von Diphthongen und Triphthongen²¹ weitaus mehr Reimmöglichkeiten als z.B. das Französische mit nur maximal 16 Vokal-phonemen.²² Sobald die den Reimvokal umgebenden Konsonanten mitreimen sollen, hängt die Anzahl der möglichen Reime auch noch von der Anzahl der Konsonanten einer Sprache und von der Anzahl der möglichen Konsonantenverbindungen ab.

Auch die Akzentuierungsgesetze einer Sprache können für die Anzahl der möglichen Reime von Bedeutung sein. Im Deutschen z.B., wo der Platz des Akzents freier ist als z.B. im Spanischen oder Italienischen, ist es schwieriger Reime zu finden, die der Bedingung der akzentuellen Identität genügen.²³

Die Abhängigkeit von der morphologischen Struktur einer Sprache zeigt sich bei einem Vergleich der Reimmöglichkeiten in einer analytischen Sprache wie dem Englischen und einer synthetischen Sprache wie dem Deutschen oder Russischen. Im Englischen gehört ein Lexem in der Regel nur einer Reimgruppe an, während im Deutschen oder Russischen ein Lexem je nach seiner Flexion verschiedenen Reimgruppen angehören kann.

Der geringe Umfang der Reimgruppen im Englischen ist auch der Grund dafür, daß das Oktett des italienischen Sonetts nicht als abba/abba, sondern als abab/cddc im Englischen realisiert wurde (vgl. LEVÝ 1965:220). Ferner sind in der englischen Lyrik sehr viele Reimwörter automatisiert. So gibt es nach LEVÝ für *love* nur drei Reimwörter (*dove, glove, above*), deren Prädiktabilität aufgrund der Information durch den Kontext praktisch 100 % ist.

Geht man mit WEAVER (1948) davon aus, daß eine im Verhältnis zu den übrigen Ereignissen eines Ereignisraumes geringe Auftretenswahrscheinlichkeit eine notwendige Bedingung für einen hohen Überraschungswert ist, ist die Wirkung eines solchen Reims natürlich nur sehr gering.

Auch die Möglichkeit der Bildung mehrsilbiger Reime hängt von der Sprachstruktur ab. Sprachen mit höherer mittlerer Silbenzahl pro Wort können bei gleichem Umfang des Lexikons und gleicher phonologischer und morphologischer Struktur eine größere Anzahl von mehrsilbigen Reimen bilden als Sprachen mit geringerer mittlerer Silbenzahl. Dabei wird die Auswahl um so kleiner, je mehr Phoneme aufeinander reimen sollen, d.h. phonetische und semantische Originalität sind in diesem Fall indirekt proportional (vgl. LEVÝ 1965:222).

8.5.2. Ein Beispiel für die einschränkende Wirkung des Zeitstils auf die Reimbildung wird von GUIRAUD (1970:87f) angeführt.

In einem französischen Reimlexikon finden sich 120 Wörter, die auf *-oire* enden. Nach GUIRAUD ergeben sich damit 15000 theoretisch mögliche Reime auf *-oire*. Interpretiert man die Zahl, die GUIRAUD ohne Begründung angibt, als $120^2 = 14400$, beruht die Berechnung der Reimmöglichkeiten bei GUIRAUD auf der Berechnung der Anzahl der Variationen mit Wiederholung. Geht man jedoch von der m.E. sinnvolleren Annahme aus, daß Reime wie $x_1 \tilde{R} y_1$ und $y_1 \tilde{R} x_1$ als gleich anzusehen sind, daß also die Anordnung der Reime unberücksichtigt bleiben sollte, reduziert sich diese Zahl auf die Anzahl der Kombinationen von 120 Elementen zur 2-ten Klasse mit Wiederholung, d.h. auf

$$C_2^w(120) = \binom{120+2-1}{2} = \frac{121!}{2! 119!} = 7260.$$

Berücksichtigt man noch, daß seit der Plejade die Reimrelation als nicht reflexiv angesehen wird, erhalten wir als theoretisch mögliche Anzahl $\binom{120}{2} = 7140$.

Eine Überprüfung des "Cid" von CORNEILLE, des "Cosroës" von ROTROU und des "Bajazet" von RACINE ergibt jedoch nur 7 verschiedene Reime auf *-oire*. Dabei machen die Reime *gloire - victoire* und *gloire - mémoire* über 50 % der Gesamtfrequenz der 7 Reime aus.

Die Gründe für die geringe Anzahl und den geringen Umfang der Reimgruppen in der Klassik liegen in der beschränkten Lexik der Literatursprache dieser Zeit, in den semantischen Kookkurrenzbe-

schränkungen und vor allem in der Verurteilung der *rime banale* (vgl. ELWERT 1966:91f).

Die Schwierigkeit, geeignete Reime zu finden, bestimmt aber auch umgekehrt die Auswahl des poetischen Vokabulars. So wächst in der französischen Poesie der Gebrauch des Wortes *ombre*, parallel zum Eindringen des Wortes *sombre* in der Sprache, stark an (vgl. GUIRAUD 1953:10).

8.5.3. Da der Reim von der Sprachstruktur abhängt, führen Sprachveränderungen - aufgrund des konservativen Charakters literarischer Tradition allerdings zuweilen erst mit einer gewissen Verzögerung - häufig auch zu einer Veränderung der Reimkonventionen.

In Frankreich haben sich z.B. im 12. Jahrhundert in mehreren Dialekten [ĩ] und [ỹ] jeweils um eine Stufe zu [ẽ] bzw. [ø] geöffnet. Dies führte dazu, daß geschriebene Assonanzen wie *i ~ in* und *u ~ un* - im Altfranzösischen waren die Nasalvokale nur schwach nasalierte kombinatorische Varianten der Oralvokale und konnten deshalb mit diesen assonieren - nicht mehr als sprachgemäß empfunden wurden. Dies dürfte mit dazu beigetragen haben, daß der Vollreim an die Stelle des Assonanzprinzips trat (vgl. GROSSE 1971:31).

9. METRIK UND SPRACHTYPOLOGIE

Die Frage nach den jeweiligen Konstituenten metrischer Systeme ist auch im Zusammenhang mit interlingualen typologischen Untersuchungen zur Metrik von Interesse. Solche Untersuchungen sind vor allem in Osteuropa durchgeführt worden. Die Zahl der bisherigen Untersuchungen¹ ist allerdings, verglichen mit der Anzahl der übrigen Publikationen in der Metrik, relativ gering.

9.1. Der m.W. umfassendste Vorschlag zu einer Typologie metrischer Systeme stammt von LOTZ (1964, 1972, 1974). Diese Typologie, die nicht nur den Anspruch auf Beobachtungsadäquatheit, sondern darüber hinaus auch auf theoretische Gültigkeit erhebt (vgl. LOTZ 1974:976), berücksichtigt allerdings nur einige wenige von LOTZ als für das Metrum notwendig erachtete Konstituenten, wie z.B. die Silbe oder die prosodischen Merkmale, nicht jedoch Erscheinungen wie Alliteration oder Reim. Wie wir gesehen haben, können jedoch auch Faktoren wie Reim, Alliteration oder syntaktisch-semantische Parallelismen in einzelnen metrischen Systemen eine bedeutende Rolle spielen.

Die Typologie von LOTZ erfüllt somit zwar möglicherweise das zur Bewertung von Typologien vorgeschlagene Kriterium der Einfachheit, nicht jedoch das Kriterium der Allgemeinheit, da lediglich einige wenige Aspekte der zu klassifizierenden Gegenstände berücksichtigt werden (vgl. ALTMANN/LEHFELDT 1973:52).

9.2. Die weitere typologische Forschung sollte zudem nicht nur primär metrische Systeme vergleichen, sondern auch verstärkt die Textebene berücksichtigen. Auf dieser Ebene müßten metrische Texte auch in bezug auf Faktoren wie Wortlänge, Phrasenlänge, Satzlänge, Satzkomplexität, Vokalhaltigkeit, Phonementropie, um nur einige wenige zu nennen, mit Hilfe quantitativer Verfahren verglichen werden. Ein solcher Vergleich ist u.a. auch deswegen von Interesse, weil diese Faktoren auch bei solchen Systemen, die LOTZ als typologisch gleich ansieht, zu beträchtlichen Unterschieden in der rhythmisch-lautlichen Gestaltung metrischer Texte führen können.²

Eine solche interlinguale Typologie setzt allerdings eine zumindest partielle einzelsprachliche Typologie metrischer Texte voraus, zumal viele Merkmale auch innerhalb eines einzelsprachlichen Texttyps, wie z.B. 'jambischer Text', erheblich variieren.

Da jedoch auch einzelsprachliche Typologien bisher lediglich in ersten Ansätzen vorliegen (z.B. für russische und tschechische bzw. slowakische Texte), ist eine allgemeine Typologie metrischer Texte höchstens als langfristiges Forschungsprogramm zu realisieren.

9.3. Eine auf einer Vielzahl von Merkmalen aufbauende allgemeine metrische Typologie, die sowohl die System- als auch Textebene berücksichtigt, sollte m.E. als Teiltheorie einer - im Sinne von ALTMANN/LEHFELDT (1973) zu verstehenden - allgemeinen Sprachtypologie konzipiert werden. Diese untersucht nicht nur die manifesten, kategorischen, sondern auch die latenten, probabilistischen Mechanismen der Sprache.

Eine so konzipierte metrische Typologie erlaubt nicht nur einen adäquateren Vergleich metrischer Texte und Systeme, sondern ist auch die Grundlage einer umfassenden linguistischen Theorie der diachronen Entwicklung metrischer Systeme.

Die Berücksichtigung auch der latenten Eigenschaften der Sprache ist für eine solche Theorie deswegen von Bedeutung, weil aufgrund der vielfältigen Beziehungen zwischen sprachlichen Einheiten auch die Veränderung eines Merkmals, das nicht als unmittelbar relevant für das metrische System angesehen wird, zu einer Systemveränderung führen kann.

So besteht z.B. ein Zusammenhang zwischen der Phonemzahl und der Wortlänge und zwischen der Wortlänge und dem Akzent. Ist der Akzent metrisch relevant, kann somit eine Veränderung der Phonemzahl bzw. der Wortlänge zumindest langfristig zu einer Veränderung des metrischen Systems führen.

Wissen wir, welche Ausprägungen eines Merkmals oder mehrerer Merkmale mit welchen Ausprägungen anderer Merkmale korreliert sind, erlaubt uns dieses Wissen, die Ausprägungen aller korrelierten

Variablen mit einer im voraus berechneten Wahrscheinlichkeit vor- auszusagen (vgl. ALTMANN/LEHFELDT 1973:58).

Dies bedeutet, daß in vielen Fällen erst eine Berücksichtigung einer Vielzahl sowohl manifester als auch latenter sprachlicher Eigenschaften Erklärungen und Prognosen in bezug auf den Wandel metrischer Systeme ermöglicht.

Allerdings dürften sich die Prognosen als außerordentlich schwierig erweisen, da der Wandel metrischer Systeme in noch weit höherem Maße als der Sprachwandel von nicht oder nur sehr schwer prognostizierbaren außersprachlichen Einflüssen abhängt.

10. DIE FUNKTION METRISCHER STRUKTUREN

Aufgrund des Schwerpunktes dieser Arbeit ist bisher die Frage der Funktion metrischer Strukturen weitgehend unberücksichtigt geblieben. Da diese Frage für eine kommunikationswissenschaftlich orientierte Metriktheorie von erheblicher Bedeutung ist (vgl. Kap. 3), sollen zumindest einige wenige Bemerkungen zu dieser Thematik erfolgen.

10.1. Entsprechend der Konzeption der Prager Strukturalisten kann die Funktion metrischer Strukturen¹ vor allem in der sog. Aktualisierung gesehen werden. Die Aktualisierung (tschech. *aktualizace*, engl. *foregrounding*) wird von HAVRÁNEK (1964:10) folgendermaßen definiert: "... the use of the devices of the language in such a way that this use itself attracts attention and is perceived as uncommon, as deprived of automatisisation...."

Der kontrastive Hintergrund für die Aktualisierung ist einmal die sog. Standardsprache und zum anderen der traditionelle ästhetische Kanon (vgl. MUKÁROVSKÝ 1964:20).

Durch die Aktualisierung wird die Sprache ihrer normalen, kommunikativen Funktion² entkleidet und die Aufmerksamkeit auf die sprachliche Form gelenkt. Diese wird zum Träger z.B. expressiver oder ästhetischer Information und zum integralen Bestandteil der Textbedeutung. Metrische Strukturen können somit eine semantisch-pragmatische Funktion erfüllen. LEVÝ (1966) spricht deshalb auch von einer "semantic theory of verse", deren Stand er folgendermaßen charakterisiert: "On the whole, the semantic theory of verse is in an unsatisfactory state of unverified, and often unverifiable conjectures." (LEVÝ 1966:139).

Dieses Urteil ist auch noch zum heutigen Zeitpunkt gerechtfertigt. Vor allem zur Frage der Wirkung metrischer Strukturen auf den Textrezipienten liegen erst relativ wenige experimentelle Untersuchungen vor.³ Von einer "Metrik als Theorie phonomorpher Wirkmuster" (SIEVEKE 1972) sind wir somit noch weit entfernt.

Die folgenden Bemerkungen müssen deshalb notwendigerweise zum Teil hypothetisch bleiben.

10.2. Die (lautliche) Form eines Textes kann zu autorspezifischen Merkmalen (z.B. biographische oder psychoanalytische Daten), zur Textbedeutung und zu Rezipientenreaktionen in Beziehung gesetzt werden.⁴ Hinsichtlich des Bezugs zur Textbedeutung und zu den Rezipientenreaktionen ergibt sich dabei vor allem das Problem, daß es oft zweifelhaft erscheint, ob die der (lautlichen) Form zugewiesene Funktion nicht nur lediglich durch die Textbedeutung suggeriert wird. Dies ist bei der Interpretation von Untersuchungen zu dieser Thematik stets zu berücksichtigen.

10.2.1. Weicht die Frequenz der Phoneme eines Textes von der durchschnittlichen Phonemfrequenz in einer Sprache ab, können lautliche Wirkungen entstehen (vgl. LEVÝ 1965:214ff). Außerdem kann auch die rhythmische Struktur und damit auch die rhythmische Wirkung eines Textes von der Phonemdistribution beeinflusst werden. Der kontrastive Hintergrund für die lautlichen Wirkungen von Phonemen ist aber nicht nur die durchschnittliche Phonemfrequenz in einer gegebenen Sprache, sondern auch die durchschnittliche Frequenz der Phoneme im entsprechenden Text.

Außerdem ist die Anordnung der Phoneme von Bedeutung. Ein Reimvokal z.B. kann aufgrund seiner exponierten Stellung weit expressiver sein als ein anderer vergleichbarer Vokal. Auch die Distanz der Phoneme muß berücksichtigt werden. Ein Vokal, der sich durch mehrere Merkmale von seiner Umgebung abhebt, dürfte von größerer Wirkung sein als ein durch ein einziges Merkmal von seiner Umgebung unterschiedener Vokal.

Wichtig ist, daß die Wirkungen in Relation zur Begrenztheit des Gedächtnisses stehen (vgl. LEVÝ 1965:215).⁵ Es werden sich vermutlich kaum irgendwelche Wirkungen ergeben, wenn sich z.B. das Phonem /f/ in einer Reihe von Lauten in jeder 25. Position befände oder nach Primzahlen geordnet wäre, obwohl dies eine deutliche Abweichung von der normalen Distribution wäre. Rekurrenzen dieser Art, wie sie zum Teil KNAUER (1965) untersucht, sind also wohl lediglich ein Indiz für das Streben des Dichters nach formaler Gestaltung, haben aber keinen expressiven Charakter. Genauere Aufschlüsse brächte aber erst eine Untersuchung der "Struktur der literarischen Wahrnehmung" und

der "maximalen und minimalen psychisch realisierbaren Komplexe" (LEVÝ 1965:215).

10.2.2. Die Frage der Korrelation zwischen lautlichen und inhaltlichen Merkmalen von Texten - häufig wird in diesem Zusammenhang von Lautsymbolik oder Lautmalerei gesprochen - ist vielfach untersucht worden.⁶ Vor allem I. FÓNAGY hat sich mit dieser Thematik - FÓNAGY spricht in Anlehnung an WUNDT (1900:326ff) von Lautmetaphern - in einer Reihe von Publikationen beschäftigt (vgl. z.B. FÓNAGY 1963).⁷ Die Ergebnisse der z.T. experimentell ausgerichteten Untersuchungen von FÓNAGY scheinen u.a. darauf hinzuweisen, daß zwischen einzelnen Lauten und damit verbundenen Vorstellungen unabhängig vom Wortschatz der einzelnen Sprachen ein "panchronischer Zusammenhang" besteht (vgl. FÓNAGY 1961:600). So herrschen z.B. im Deutschen, Französischen und Ungarischen in aggressiven Gedichten harte Laute (r,k,t), in zärtlichen Gedichten dagegen weiche Laute (l,m) vor (vgl. FÓNAGY 1965:245).

Eine Korrelation zwischen der Lautgestalt und dem Inhalt sprachlicher Zeichen findet sich jedoch nicht nur in der Poesie. Wie vor allem Psychologen experimentell nachgewiesen haben,⁸ ist die Lautgestalt sprachlicher Zeichen keineswegs so arbiträr wie von Linguisten häufig angenommen wird.⁹ Deutlich wird dies von dem Psychologen ERTEL, der in seiner Habilitationsschrift zu dieser Frage zahlreiche Tests durchgeführt hat, ausgesprochen: "Das Prinzip der völligen 'Arbitrarität' oder Beliebigkeit phonologisch semantischer Relationen widerspricht den Tatsachen." (ERTEL 1969:199)

Ein Beispiel für die Nicht-Arbitrarität phonologisch-semantischer Relationen ist das in einer Reihe von Sprachen nachgewiesene Phänomen, daß in Wörtern, die 'Helligkeit' bezeichnen, die Frequenz der hellen Vokale eindeutig höher ist, als in Wörtern, die 'Dunkelheit' bezeichnen (vgl. z.B. CHASTAING 1964; FÓNAGY 1971). Diese latenten phonologisch-semantischen Relationen der Sprache können bei der Textgestaltung

ebenfalls als stilistisches Mittel eingesetzt werden (vgl. FÓNAGY 1971). Daß sich die Dichter dieser Möglichkeit durchaus bewußt sind, zeigt die folgende Bemerkung MALLARMÉs in den "Variations sur un sujet": "Quelle déception, devant la perversité conférant à j o u r comme à n u i t, contradictoirement, des timbres obscur ici, là clair."¹⁰

10.2.3. Von den Fällen, bei denen möglicherweise bereits a priori ein Zusammenhang zwischen der Lautung und der Bedeutung besteht, ist die für die Dichtung typische sekundäre Motivation zu trennen, bei der erst a posteriori die Lautung durch die Bedeutung motiviert wird.¹¹

Ein interessantes Beispiel findet sich bei ŠABRŠULA (1969:182f). In einem Vers wie

Et le peuple en rumeur gronde autour du prétoire
(LECONTE DE LISLE, La passion)

stellen wir mehrere /r/ fest, die ein dem Wort 'gronde' analoges Geräusch ausdrücken. In einem Vers wie

Ton corps rare, harmonieux
(VERLAINE)

drückt der gleiche Konsonant /r/ jedoch nicht das gleiche aus wie in den obigen Versen von LECONTE DE LISLE.

Auch in Wörtern wie z.B. *rime*, *radieux* drückt das Phonem /r/ auf keinen Fall ein vergleichbares Geräusch aus.

Der Rhythmus kann ebenfalls sekundär motiviert sein. So scheint z.B. in dem Vers

Infinis bercements du loisir embaumé
(BAUDELAIRE, La Chevelure)

durch die viermalige Wiederholung des Schemas xx̣ der Rhythmus eine wiegende Bewegung (*bercement*) auszudrücken. In dem Vers

Ces serments, ces parfums, ces baisers infinis
(BAUDELAIRE, Le Balcon)

drückt der gleiche Rhythmus jedoch keine solche Bewegung aus.¹²

Diese und eine Vielzahl weiterer Beispiele zeigen, daß Phonemdistribution und Rhythmus meist keine eindeutig bestimmbare, autonome kognitive Bedeutung haben (vgl. CHATMAN 1965:198; LOTMAN 1973:242f). Dies dürfte auch für die übrigen formalen Texteigenschaften gelten.

10.2.4. Auch auf der syntaktischen Ebene gibt es die Möglichkeit, Verse expressiv zu gestalten. So sind z.B. in den lebhaften Satiren des HORAZ, die meist alltägliche Dinge zum Thema haben, in Analogie zur Alltagssprache, wo oft in Erregung ein Laut oder eine Silbe verschluckt wird, Aphärese und Elision weit häufiger als in den Briefen von ruhigem Ton (vgl. NILSSON 1952; FÓNAGY (1965:245)). Auch Zäsur und Zeilenbruch können, wie FÓNAGY (1965:245f) anhand von statistischen Untersuchungen an deutschen, französischen und ungarischen Dichtern festgestellt hat, sowohl mit inhaltlichen als auch mit autorspezifischen Persönlichkeitsmerkmalen korrelieren (vgl. auch FÓNAGY 1971a).

Auch die Versart kann ein Hinweis auf eine bestimmte Art des Inhalts sein. LOPE DE VEGA in seiner "Arte nuevo de hacer comedias" gibt folgende Empfehlung:¹³

Acomode los versos con prudencia
a los sujetos de que va tratando

Dann folgt eine Aufzählung, welche Versarten zu welchem Inhalt passen, z.B.

Las décimas son buenas para quejas;

Mit dieser Empfehlung hat LOPE DE VEGA aber lediglich die herrschende Meinung der Poetiken in Verse gebracht. Denn schon in der Antike gelten bestimmte Versarten an einen bestimmten Inhalt gebunden. So ist im Lateinischen der Hexameter der Epenvers, und das Distichon gilt als der elegische Vers par excellence. Jambische Verse dagegen gelten aufgrund ihrer großen Modulationsfähigkeit als verschiedenen Inhalten angemessen (vgl. KAYSER 1967:257ff).

Abschließend sei noch auf ein Experiment hingewiesen, das für die Beurteilung von Arbeiten zur Funktion der lautlichen Form von Bedeutung ist.

JACOB (1918:55) hat einer Gruppe von Studenten, die kein Deutsch konnten, stark lautmalende deutsche Gedichte vorlesen lassen. Es stellte sich heraus, daß der auf die Studenten gemachte Eindruck von der Vortragsart abhing (lustig, pastoral etc.) und nicht von der Vokal- und Konsonantenqualität.

Da in vielen Arbeiten zur lautlichen Wirkung bzw. zur Lautsymbolik kaum deutlich wird, ob und in welcher Form der Einfluß der Variablen 'Inhalt' und 'Vortragsart' kontrolliert worden ist, dürfte der Ausgang des Experiments von JACOB ein (weiterer) Grund sein, die Validität der Ergebnisse dieser Arbeiten in Zweifel zu ziehen.

11. ZUR THEORIE DER MATHEMATISCHEN METRIK

11.1. Die Verwendung mathematischer Methoden in der Metrik

11.1.1. Die Entwicklung der empirischen Wissenschaften führt - und dies gilt nicht nur für die Naturwissenschaften, sondern auch, wie das Beispiel von Disziplinen wie Psychologie oder empirische Sozialforschung zeigt, für die sog. Humanwissenschaften - notwendigerweise zu einem Punkt, wo erst durch die Verwendung mathematischer Methoden eine tiefergehende Erkenntnis des Gegenstandsreiches ermöglicht wird (vgl. z.B. FREY 1967; BOUDON 1972). Die Gründe hierfür sind vielfältiger Natur. Ein wichtiger Grund ist, daß die Mathematik eine kürzere und präzisere Darstellung und Argumentation ermöglicht als nicht-mathematische Sprachsysteme. Daneben bietet vor allem die stochastische Mathematik eine Reihe von weiteren Vorteilen (s.u.).

Es ist somit wenig erstaunlich, daß auch in den sprachthematizierenden Wissenschaften eine immer stärker werdende Tendenz zur Verwendung mathematischer Verfahren zu beobachten ist. Diese Entwicklung, die vor allem von Mathematikern und Linguisten ausgegangen ist, hat mittlerweile zur Etablierung einer mathematischen Linguistik und zur Herausbildung einer Reihe sich z.T. überschneidender und auch z.T. unterschiedlich bezeichneter Forschungsrichtungen geführt, die sich speziell mit der mathematischen Analyse von (literarischen) Texten beschäftigen. Als umfassendere Richtungen sind hier vor allem die mathematische Texttheorie (vgl. FISCHER 1976), die mathematische Poetik (vgl. MARCUS 1973) und die statistische Stilistik (vgl. DOLEŽEL/BAILEY 1969) zu nennen, die - allerdings mit unterschiedlichen Erkenntnisinteressen - auch die mathematische Analyse metrischer Texte zu ihren Forschungszielen rechnen.

11.1.2. In der mathematischen Metrik¹ können ebenso wie in der mathematischen Linguistik (vgl. ALTMANN 1972; 1973) oder in der mathematischen Texttheorie (vgl. FISCHER 1976) zwei Hauptrichtungen unterschieden werden:

a) Eine primär algebraisch orientierte Richtung, die als *algebraische Metrik* bezeichnet werden soll. Diese untersucht die qualitativen Eigenschaften metrischer Texte und verwendet z.B. folgende mathematische Verfahren: Logik, Mengentheorie, Gruppentheorie, Verbandstheorie, Graphentheorie, Topologie, Geometrie. Am bekanntesten sind wohl in diesem Bereich die Arbeiten des rumänischen Mathematikers S. MARCUS (vgl. z.B. MARCUS 1973).

b) Eine primär quantitativ orientierte Richtung, die als *quantitative Metrik* bezeichnet werden soll. Diese quantifiziert die Eigenschaften metrischer Texte, d.h. schreibt diesen Eigenschaften Maße zu. Mit Hilfe wahrscheinlichkeitstheoretischer und statistischer Verfahren versucht sie auch dort Regelmäßigkeiten und Zusammenhänge aufzudecken, wo algebraische Methoden nicht mehr weiterführen. Während die algebraischen Verfahren vor allem die kategorischen und manifesten Eigenschaften erfassen,² können mit Hilfe quantitativer (statistischer) Verfahren auch *latente*, gleichsam unter der Textoberfläche liegende *Tendenzen* erfaßt werden. In der quantitativen Metrik werden z.B. folgende mathematische Verfahren verwendet: Analysis, Kombinatorik, Wahrscheinlichkeitstheorie (einschließlich stochastischer Prozesse und Informationstheorie), Statistik, numerische Taxonomie, numerische Mathematik.

In der mathematischen Linguistik hat sich zwar insbesondere im Bereich der Lexikologie (vgl. z.B. MULLER 1972) und der Phonetik (vgl. z.B. ALTMANN 1971a; ALTMANN/LEHFELDT 1979) eine quantitativ orientierte Richtung etabliert, doch stand vor allem im Zusammenhang mit der Entwicklung formaler Grammatiken die algebraische Richtung lange Zeit im Vordergrund des Forschungsinteresses. In jüngster Zeit scheint jedoch vor allem im Zusammenhang mit der Untersuchung sprachlicher Variation (auch bei algebraisch orientierten Linguisten) die Bereitschaft zu wachsen, durch die Verwendung statistischer Verfahren dem stochastischen Charakter der Sprache Rechnung zu tragen (vgl. z.B. KLEIN 1974a). Von einigen Autoren wird diese Entwicklung sogar als Beginn eines Paradigmenwechsels angesehen (vgl. CEDERGREN/SANKOFF 1974:334).

Die mathematische Metrik ist dagegen bis zum heutigen Tage vor allem quantitativ orientiert. Der Grund hierfür dürfte vor allem darin zu sehen sein, daß in mathematischen Untersuchungen zur Metrik in der Regel der konkrete Text als künstlerisches, in seiner Gestaltung von einer Vielzahl von Faktoren abhängiges Produkt eines individuellen Textproduzenten im Vordergrund der Betrachtungsweise steht. Metrische Phänomene werden vorwiegend als Stilphänomene angesehen und 'Stil' läßt sich - diese Einsicht hat sich in den letzten Jahrzehnten immer mehr durchgesetzt - adäquat nur mit quantitativen (statistischen) Methoden erfassen.

Quantitative Verfahren sind aber nicht nur bei der Stilanalyse von Bedeutung. Wir hatten in den Kapiteln 6 und 9 gesehen, daß auch z.B. die Untersuchung der Abhängigkeit metrischer Systeme vom Sprachsystem oder die Typologie metrischer Systeme und Texte von großer Bedeutung für die Metriktheorie ist. Auch hier ist die Verwendung quantitativer Verfahren eine wichtige Voraussetzung für die Weiterentwicklung dieser Forschungsrichtungen.

Mit der Betonung der Bedeutung quantitativer Verfahren soll jedoch nicht die Notwendigkeit der Verwendung algebraischer Verfahren verneint werden. Auch in der Metrik gibt es Zusammenhänge und Eigenschaften, die sinnvoll mit Hilfe algebraischer Methoden erfaßt werden können (vgl. MARCUS 1973 und Kap. 4 und 8 dieser Arbeit). Algebraische und quantitative Metrik sind somit nicht als konkurrierende Forschungsrichtungen anzusehen, sondern verfolgen dasselbe Ziel, nämlich eine tiefere Einsicht in den Gegenstandsbereich der Metrik zu gewinnen.

Entsprechend dieser Konzeption kann die mathematische Metrik wissenschaftssystematisch auch am ehesten einer sowohl schriftliche als auch mündliche Texte untersuchenden mathematischen Texttheorie zugeordnet werden.³

11.1.3. Die Anfänge der systematischen Verwendung quantitativer Methoden bei der Analyse metrischer Texte gehen vor allem auf die russischen Formalisten und die tschechischen und polnischen Strukturalisten zurück.

Dies ist wenig überraschend, denn gerade bei diesen Schulen spielt die Analyse der Formalstruktur literarischer Texte eine zentrale Rolle (vgl. LEVÝ 1965, DOLEŽEL 1967).

Nach LEVÝ (1965) beginnt die Entwicklung der mathematischen Versanalyse 1910 mit dem Werk von BELYJ über den Symbolismus (vgl. auch LOTMAN 1973:178). Diese Datierung ist insofern richtig, als hier zum ersten Mal ein systematisches Programm für die statistische Analyse metrischer Strukturen aufgestellt worden ist. Aber bereits mehr als 40 Jahre früher wurden von DROBISCH (1866) bei der Analyse des lateinischen Hexameters statistische Verfahren verwendet.

Im Anschluß an BELYJ haben sich in Osteuropa zahlreiche Forscher mit Fragen der quantitativen Metrik beschäftigt, so z.B. in Rußland die Gruppe um den Begründer der axiomatischen Wahrscheinlichkeitstheorie A.N. KOLMOGOROV, in Polen J. WORONCZAK und die Gruppe um M.R. MAYENOWA und in der Tschechoslowakei J. LEVÝ und M. ČERVENKA u.a. Dagegen sind außerhalb Osteuropas nur relativ wenige Arbeiten zur quantitativen Metrik zu verzeichnen.⁴

Da ein umfassender Überblick über die Methoden und Ergebnisse der mathematischen Metrik eine gesonderte Darstellung mit bedeutendem Umfang erfordern würde, werde ich im Rahmen dieser Arbeit nur einige Methoden der quantitativen Analyse metrischer Texte exemplarisch aufzeigen können. Dies soll in Kapitel 12 geschehen. Zuvor sollen jedoch noch einige für die quantitative Textanalyse - und damit auch für das Kapitel 12 - wichtige theoretische Probleme diskutiert werden.

11.2. Ein Begriffsrahmen für die quantitative Textanalyse

11.2.1. Ein Text ist das Produkt eines komplexen Entscheidungsprozesses über die Auswahl und Anordnung des sprachlichen Materials. Die Entscheidungen des Textproduzenten - und damit auch die daraus resultierende spezifische Textstruktur - sind sowohl vom jeweiligen Sprachsystem als auch von einer Vielzahl weiterer Faktoren abhängig, die global als pragmatische Faktoren bezeichnet werden können (vgl. Kap. 3 und DOLEŽEL 1969).⁵

L. DOLEŽEL unterscheidet zwei Hauptgruppen von pragmatischen Faktoren: die subjektiven und die objektiven pragmatischen Faktoren. Zu den subjektiven pragmatischen Faktoren rechnet er die psycholinguistischen Eigenschaften des Textproduzenten, wie Niveau der Sprachbeherrschung, verbale Präferenzen, mentaler Typ, ständiger und momentaner psychischer Zustand usw.

Zu den objektiven pragmatischen Faktoren zählt DOLEŽEL die auf den Textproduzenten einwirkenden Faktoren, wie Funktion des Textes, Genre, Übertragungskanal (schriftlich/mündlich), poetischer Kanon usw. Diese Faktoren werden von ihm unter dem Begriff des Kontexts zusammengefaßt (DOLEŽEL 1965; 1969).

Bei der Textanalyse ist zu unterscheiden, ob eine Texteigenschaft allein vom Sprachsystem abhängt oder auch von pragmatischen Faktoren. Wenn die numerischen Werte einer Texteigenschaft - die numerischen Werte einer Eigenschaft sollen im folgenden Charakteristiken (C) genannt werden - in der Menge der Texte einer Sprache L zu einem Zeitpunkt t stabil sind, d.h. relativ zu einem bestimmten statistischen Zufallsmodell keine signifikanten Differenzen aufweisen, ist anzunehmen, daß sie allein von der Sprache L abhängen und nicht von pragmatischen Faktoren. Solche Charakteristiken nennt DOLEŽEL (1969:18) L-Charakteristiken (L-C).

L-Charakteristiken können nicht als Textdiskriminatoren benutzt werden und sind deshalb als suprastilistische Charakteristiken anzusehen. Ein Beispiel für eine L-Charakteristik ist die Anzahl der verschiedenen Phoneme in einem Text, die in der Regel nicht von pragmatischen Faktoren abhängt.

Variiert dagegen eine Textcharakteristik in Abhängigkeit von pragmatischen Faktoren, handelt es sich um eine stilistische Textcharakteristik, im folgenden Stilcharakteristik (SC) genannt (vgl. DOLEŽEL 1969). Ob eine solche Abhängigkeit vorliegt, kann durch eine systematische Variation der unabhängigen Variablen 'Textproduzent' und/oder 'Kontext' überprüft werden. Führt eine solche Variation zu statistisch signifikanten Differenzen zwischen den abhängigen Variablen 'Textcharakteristiken' ist dies, sofern der Einfluß weiterer Variablen entweder experimentell oder durch entsprechende statistische Verfahren, wie z.B. Kovarianzanalyse, kontrolliert worden ist, ein Hinweis auf einen Kausalzusammenhang zwischen den genannten Variablen.⁶

Stilcharakteristiken können im Gegensatz zu L-Charakteristiken zur Textdifferenzierung benutzt werden. Ein Beispiel für eine Stilcharakteristik ist die mittlere Wortlänge oder die mittlere Satzlänge, die bei verschiedenen Autoren und/oder in verschiedenen Kontexten instabil sind, d.h. statistisch signifikante Unterschiede aufweisen.

Durch die Operationalisierung des Begriffs 'Stilcharakteristik' ist es möglich, auch den Begriff des Stils zu operationalisieren.

Da Stilcharakteristiken die Ausprägung von Stileigenschaften eines Textes bzw. einer Klasse von Texten kennzeichnen,⁷ kann unter 'Stil' die Gesamtheit der Stilcharakteristiken eines Textes oder einer Klasse von Texten verstanden werden (vgl. DOLEŽEL 1969).⁸ Diese Sicht des Phänomens Stil hat zur Folge, daß es nicht als ausreichend angesehen werden kann, wenn - wie es häufig in der Stilforschung der Fall ist - zur Kennzeichnung des Stils eines Textes nur eine einzige Charakteristik benutzt wird. Eine einzelne Charakteristik mit hoher Variabilität kann zwar ein guter Stildiskriminator sein, für eine vollständige Charakterisierung des Stils eines Textes müssen jedoch auch Stilcharakteristiken mit geringer Variabilität berücksichtigt werden. Denn auch diese kennzeichnen, sofern ihre Variabilität nicht dem Zufall zuzuschreiben ist, den Stil eines Textes.

11.2.2. Metrische Strukturen haben grundsätzlich stilistischen Charakter, da sie nicht in allen Texten einer Sprache L zu einem Zeitpunkt t stabil sind. Es ist somit angebracht, sich eingehender mit der Frage der Stilcharakteristiken zu beschäftigen.

Es kann mit DOLEŽEL (1969) zwischen subjektiven, objektiven und subjektiv-objektiven Stilcharakteristiken unterschieden werden.

- (1) Eine Stilcharakteristik heißt subjektiv (S-C), wenn sie von dem Faktor 'Textproduzent' abhängt, aber nicht von dem Faktor 'Kontext'.
- (2) Eine Stilcharakteristik heißt objektiv (O-C), wenn sie von dem Faktor 'Kontext' abhängt, aber nicht von dem Faktor 'Textproduzent'.

- (3) Eine Stilcharakteristik heißt s u b j e k t i v - o b j e k t i v (S-O-C), wenn sie sowohl vom Faktor 'Textproduzent' als auch vom Faktor 'Kontext' abhängt.

DOLEŽEL schlägt vor, die S-O-Charakteristiken aus den S-Charakteristiken mit Hilfe des folgenden Testverfahrens abzuleiten:

"(1) Let us test the hypothesis that the fluctuation of $(S-C_i)$ in the corpus $T(Q_j)$ falls within the interval $\langle a; k \rangle$, whereas its fluctuation in the corpus $T(Q_k)$ is confined to the interval $\langle l; z \rangle$. The two intervals are, of course, not overlapping.

(2) Let us test the hypothesis that the average value (mathematical expectation) of $(S-C_i)$ in $T(Q_j)$ is significantly different from its average value in $T(Q_k)$.

If the result of either of these tests is positive, we can conclude that the $S-C_i$ of the corpus $T(Q_j)$ fluctuates within a certain range which is different from that of another corpus, $T(Q_k)$. The tested characteristic thus apparently expresses both a subjective style feature (by its significant fluctuation within the corpus) and an objective style feature (by the demarcated range of fluctuation)." (DOLEŽEL 1969:20f)⁹

Die beiden Testverfahren können jedoch zu falschen Ergebnissen führen, da sie nicht die Möglichkeit einer Wechselwirkung (Interaktion) zwischen den Faktoren 'Textproduzent (TP)' und 'Kontext (K)' berücksichtigen. Dies kann leicht anhand eines hypothetischen Beispiels verdeutlicht werden.

Angenommen wir erhalten bei der Analyse der Satzlänge (gemessen in Wörtern) folgende Mittelwerte:

	K ₁	K ₂	
TP ₁	20	40	30
TP ₂	40	20	30
	30	30	

Beide von DOLEŽEL vorgeschlagenen Testverfahren führen in einem solchen Fall trotz des offensichtlichen Einflusses des Kontextes zu einem nicht signifikanten Ergebnis. Dies kann jedoch leicht vermieden werden, wenn man von einem faktoriellen Versuchsplan ausgeht und z.B. eine faktorielle Varianzanalyse durchführt.

Ein solches Vorgehen hat darüber hinaus noch einen weiteren Vorteil. Es ist durchaus möglich, daß der Einfluß der Faktoren 'Textproduzent' und 'Kontext' einzeln betrachtet nicht signifikant ist, daß die beiden Faktoren jedoch zusammen einen signifikanten Einfluß auf eine Stilcharakteristik ausüben.

Mit Hilfe der Merkmale 'subjektiv' und 'objektiv' kann somit zwischen insgesamt vier verschiedenen Typen von Stilcharakteristiken differenziert werden:

	Stilcharakteristik			
	C ₁	C ₂	C ₃	C ₄
subjektiv	+	-	+	-
objektiv	-	+	+	-

Der Typ C₄, der von DOLEŽEL nicht berücksichtigt wird, könnte als nichtsubjektiv-nichtobjektive Stilcharakteristik (NS-NO-C) bezeichnet werden (vgl. ALTMANN 1971).

Formal kann der Stil eines Textes jetzt als Menge von subjektiven, objektiven, subjektiv-objektiven und nichtsubjektiv-nichtobjektiven Stilcharakteristiken dargestellt werden:

$$S = \{ S-C_1, \dots, S-C_k, \quad O-C_{k+1}, \dots, O-C_m, \\ S-O-C_{m+1}, \dots, S-O-C_n, \quad NS-NO-C_{n+1}, \dots, NS-NO-C_r \}$$

Berücksichtigt man noch die Korrelation der Stilcharakteristiken untereinander, kann S durch Ausschluß redundanter Charakteristiken auf eine minimale Anzahl grundlegender Charakteristiken reduziert werden, die vollkommen voneinander unabhängig sind und von denen man

alle anderen ableiten kann (vgl. DOLEŽEL 1969:19f; ALTMANN 1971: 278).

Durch den Typ und die Ausprägung der verwendeten Stilcharakteristiken wird die Individualität und die kontextuelle Bedingtheit eines Stils in exakter Weise charakterisiert. Die Auflösung des Stils in eine Menge von distinktiven Stilcharakteristiken ermöglicht somit auch eine exaktere und objektivere Bestimmung traditioneller Stil kategorien wie Individualstil oder Gattungsstil.

11.2.3. Bisher habe ich 'Stil' als statisches Phänomen betrachtet. Bei der Stilanalyse ist jedoch auch zu berücksichtigen, daß Stilcharakteristiken vom Faktor 'Zeit' abhängen können. So kann sich z.B. der Stil eines Autors im Laufe der Zeit verändern, oder eine Stilcharakteristik kann im Verlauf eines Textes unterschiedliche Werte annehmen.

Dieser Tatsache kann durch die Unterscheidung von *stationären* und *nicht-stationären* Stilcharakteristiken Rechnung getragen werden (vgl. DOLEŽEL 1965, 1967).¹⁰ Für eine stationäre Stilcharakteristik gilt, daß ihre Werte von der Variablen 'Zeit' unabhängig sind. Ist die Bedingung der Unabhängigkeit nicht erfüllt, handelt es sich um eine nicht-stationäre Stilcharakteristik. Nicht-stationäre Stilcharakteristiken kennzeichnen die Dynamik des Stils. Sie lassen sich als mathematische Funktionen darstellen, deren Parameter mit Hilfe regressionsanalytischer Methoden bestimmt werden können.

Es können drei Typen von stationären bzw. nicht-stationären Stilcharakteristiken unterschieden werden:

- (1) Eine Stilcharakteristik kann die Veränderung eines Individualstils (Jugendstil - Altersstil) oder eines Genrestils in Abhängigkeit vom Faktor 'Zeit' kennzeichnen. Eine solche nicht-stationäre Charakteristik soll als *diachrone* Stilcharakteristik bezeichnet werden. Bei der Untersuchung diachroner Stilcharakteristiken ist zu berücksichtigen, daß die Nicht-Stationarität auch durch die diachrone Entwicklung

(Nicht-Stationarität) des Sprachsystems bedingt sein kann.

- (2) Eine Charakteristik kann in ihrem Werteverlauf von der Zunahme des Textumfanges abhängen. Ein Beispiel für eine solche Charakteristik ist die Type-Token-Relation lexikalischer Einheiten. Berechnet man den Wert dieser Charakteristik zuerst für die ersten k Einheiten eines Textes, dann für immer größer werdende Teilmengen und schließlich für den Gesamttext, ergibt sich eine hohe Korrelation mit der Variablen 'Textumfang'. Eine nicht-stationäre Stilcharakteristik dieses Typs soll mit MÜLLER (1972:164) als *evolutorische* Stilcharakteristik bezeichnet werden. Es ist anzunehmen, daß evolutorische Charakteristiken vor allem vom Sprachsystem abhängen und nur wenig dem Einfluß pragmatischer Faktoren unterliegen.
- (3) Eine Charakteristik kann im Verlauf eines Textes z.B. in Abhängigkeit von der Akteinteilung eines Schauspiels oder von der Stropheneinteilung eines Gedichtes unterschiedliche Werte annehmen. Eine nicht-stationäre Stilcharakteristik dieses Typs soll als *intratextuelle* Stilcharakteristik bezeichnet werden.

Die Merkmale intratextuell, evolutorisch und diachron können auch kombiniert auftreten. So ist es z.B. durchaus denkbar, daß sich eine evolutorische Stilcharakteristik in Abhängigkeit vom Alter eines Autors verändert. Außerdem können alle drei Merkmale mit den Merkmalen subjektiv, objektiv, subjektiv-objektiv und nichtsubjektiv-nichtobjektiv kombiniert werden. Auf diese Weise lassen sich bereits eine Reihe verschiedener Typen von Stilcharakteristiken unterscheiden, die wiederum zur Differenzierung zwischen verschiedenen Stiltypen benutzt werden können.

11.3. Der forschungslogische Ablauf quantitativer Textanalysen

Für die quantitative Textanalyse gilt wie für jede quantitativ ausgerichtete empirische Untersuchung, daß ein bestimmter forschungslogischer Ablauf eingehalten werden muß, wenn man zu adäquaten Ergebnissen gelangen will. Es können insgesamt 5 Schritte unterschieden werden (vgl. zum Folgenden vor allem ALTMANN 1973; ferner FRIEDRICHS 1973; FRIEDRICH/HENNIG 1975).

(1) *Hypothesenbildung*. Der erste wichtige Schritt besteht in der Formulierung von Hypothesen, die durch die Untersuchung überprüft werden sollen.¹¹ Solche Hypothesen, die auch als Forschung- oder Arbeitshypothesen bezeichnet werden, beruhen auf bestimmten Annahmen über die Beschaffenheit des jeweiligen Forschungsgegenstandes und haben sprachlich die Form eines Aussagesatzes (vgl. KERLINGER 1973:18ff). Hypothesen müssen einer Reihe von Bedingungen genügen (vgl. hierzu z.B. KORCH u.a. 1972). Die wichtigste Bedingung ist die Forderung nach empirischer Überprüfbarkeit (vgl. Kap. 2). Die empirische Überprüfbarkeit einer Hypothese steht in engem Zusammenhang mit dem jeweiligen Hypothesentyp.¹² So kann z.B. eine unbeschränkte Allaussage lediglich falsifiziert, nicht aber verifiziert werden. Unbeschränkte Existenzaussagen können dagegen lediglich verifiziert, nicht jedoch falsifiziert werden. Singuläre Aussagen sind sowohl verifizierbar als auch falsifizierbar.

Im Zusammenhang mit der quantitativen Textanalyse könnten z.B. folgende Hypothesen formuliert werden: "Die Autoren x und y verwenden unterschiedlich lange Sätze", "Es besteht ein Unterschied im Verhältnis der Anzahl der Adjektive zur Anzahl der Verben zwischen den frühen und den späteren Schriften GOETHEs", "Im lateinischen Hexameter besteht eine Tendenz, am Versanfang vor allem Daktylen zu verwenden" usw.

(2) *Formalisierung*. Wenn eine Hypothese formuliert worden ist, muß sie in die Sprache eines mathematischen Modells übersetzt werden. Zu diesem Zweck müssen die verwendeten Begriffe u.a. operationalisiert werden, d.h. es müssen präzise Regeln angegeben werden, mit deren Hilfe eindeutig entscheidbar ist, ob ein bestimmtes empirisches Phänomen unter den entsprechenden Begriff fällt oder nicht. Handelt es sich um einen Begriff mit indirektem empirischen Bezug, geht der Operationalisierung die Bildung von Indikatoren voraus (vgl. MAYNTZ/HOLM/HÜBNER 1972:18ff). Welche Indikatoren und welche Operationalisierung gewählt wird, hängt vom jeweiligen Untersuchungsziel ab. So kann es z.B. angebracht sein, für eine typologische Untersuchung eine andere Operationalisierung des Silbenbegriffes zu wählen als für eine Untersuchung im Bereich der Metrik (vgl. Kap. 7).

Die Adäquatheit der Indikatorenbildung und der Operationalisierung wird anhand der Kriterien der Validität (Gültigkeit)¹³ und der Reliabilität (Zuverlässigkeit) überprüft. Indikatoren und Operationalisierung sind dann valide, wenn sie das erfassen, was sie erfassen sollen. Sie sind reliabel, wenn sie das, was sie erfassen, genau erfassen. Reliabilität ist eine notwendige, jedoch keine hinreichende Bedingung für die Validität. Das bedeutet z.B., daß zwar ohne eine exakte Operationalisierung keine gültigen Ergebnisse erzielt werden können, daß aber auch eine noch so exakte Operationalisierung nicht notwendigerweise zu validen Begriffen führt.¹⁴

Nach erfolgter Indikatorenbildung und Operationalisierung ist zu entscheiden, welches mathematische Modell für das Forschungsproblem am adäquatesten ist, d.h. welches mathematische Modell die als relevant erachteten Eigenschaften und Relationen des Modelloriginals am genauesten erfaßt (vgl. STACHOWIAK 1973).

Gerade bei der statistischen Textanalyse wird in diesem Zusammenhang häufig vergessen, daß jede Anwendung der Statistik auf bestimmten Modellvorstellungen von z.B. unabhängigen Faktoren, linearen Zusammenhängen, bestimmten Verteilungsfunktionen, Skaleneigenschaften usw. beruht und daß unterschiedliche Modellvorstellungen bei den-

selben Daten zu völlig anderen 'Ergebnissen' führen können (vgl. KRIZ 1973:24).

Bei der statistischen Textanalyse kommt es auch relativ häufig vor, daß kein adäquates statistisches Modell vorhanden ist. In einem solchen Fall muß der Forscher selbst ein statistisches Modell entwickeln. Dies setzt allerdings bereits eine gründliche Vertrautheit mit den mathematischen Grundlagen der Statistik voraus.

(3) *Datenerhebung und Datenanalyse.* Wenn die Hypothese in die Sprache des statistischen Modells übersetzt worden ist, muß zur Prüfung der Hypothese eine große Menge von Daten erhoben werden. Dabei sollten soweit wie möglich die Prinzipien der mathematischen Stichprobentheorie beachtet werden (vgl. hierzu z.B. COCHRAN 1963; MENGES/SKALA 1973). Dies ist insbesondere dann angebracht, wenn aufgrund der Stichprobenwerte Aussagen über die Grundgesamtheit gemacht werden sollen. Vor allem ist jedoch wichtig, daß die Daten zur Prüfung und nicht nur zur Rechtfertigung der eigenen Untersuchungshypothesen erhoben werden, wie dies häufig in der Linguistik, insbesondere jedoch in der Literaturwissenschaft der Fall ist.

Die gewonnenen Daten bilden das Ausgangsmaterial für die statistische Analyse. Hierbei werden u.a. Stichprobencharakteristiken berechnet, Parameter geschätzt und Signifikanztests durchgeführt. In dieser Forschungsphase ist vor allem bei größeren Datenmengen und komplizierteren statistischen Analysen der Computer ein wichtiges Hilfsmittel.

(4) *Entscheidung.* Die statistische Analyse der Daten bildet die Grundlage für die Entscheidung über Annahme oder Ablehnung einer in der Sprache des statistischen Modells formulierten Hypothese. Eine solche Entscheidung wird fast stets aufgrund von Stichprobenwerten getroffen und beruht schon deshalb notwendigerweise auf unvollständiger Information über den jeweiligen Untersuchungsgegenstand. Außerdem kennen wir in der Regel nur einen Teil der Faktoren, die das Verhalten des Untersuchungsgegenstandes bestimmen. Erst die Inferenzstatistik eröffnet in diesem Fall die Möglichkeit, trotz unvollständiger Information zu einer in bezug auf die jeweilige Informations-

lage optimalen, intersubjektiv überprüfbaren Entscheidung zu gelangen (vgl. RIEGER 1972; WICKMANN 1972). Eine quantitative Textanalyse ohne Inferenzstatistik, wie sie leider immer noch anzutreffen ist, bleibt dagegen weitgehend spekulativ und ist deshalb als Grundlage für eine Entscheidung über die Validität einer Hypothese von äußerst zweifelhaftem Wert.

Die Statistik ist somit nicht nur aus 'ontologischen' Gründen (stochastischer Charakter des Gegenstandsbereiches), sondern auch aus erkenntnistheoretischen Gründen ein unentbehrliches Werkzeug der Textanalyse.

(5) *Interpretation.* In einem letzten Schritt muß die getroffene Entscheidung interpretiert werden. Dazu muß sie zuerst in die Sprache der Ausgangshypothese zurückübersetzt werden. Anschließend ist zu prüfen, welche Konsequenzen aus dem Ergebnis der Hypothesenprüfung zu ziehen sind. Dies kann z.B. dazu führen, daß die Ausgangshypothese präzisiert oder erweitert oder auch durch eine sinnvollere Hypothese ersetzt wird. Ist die Hypothese Teil eines Hypothesensystems, muß geprüft werden, welche Auswirkungen das Ergebnis auf die übrigen Hypothesen des Systems hat. Falls die Hypothese im Zusammenhang mit einer literaturwissenschaftlichen Fragestellung steht, ist das Ergebnis auch auf diesem Hintergrund zu interpretieren. Steht die Hypothese im Widerspruch zu einer Theorie, ist zu entscheiden, ob diese Theorie zugunsten einer Alternativtheorie - sofern eine solche vorhanden ist - verworfen werden soll oder ob z.B. von der Möglichkeit der Exhaustion, d.h. der Stützung der Theorie durch Ad-hoc-Hypothesen Gebrauch gemacht werden soll (vgl. Kap. 2).

Sowohl die Phase der Interpretation als auch die Phase der Hypothesenbildung setzen eine gründliche Kenntnis der Textwissenschaft (Linguistik/Literaturwissenschaft) voraus. Dagegen ist in den übrigen Phasen eine gründliche Vertrautheit mit der Statistik eine unabdingbare Voraussetzung. In den wenigsten Fällen ist jedoch der Textwissenschaftler gleichzeitig auch ein guter Statistiker. Umgekehrt sind die wenigsten Statistiker auch gleichzeitig gute Text-

wissenschaftler.¹⁵ Dies hat bisher dazu geführt, daß eine Reihe von Textwissenschaftlern ihre Hypothesen entweder gar nicht oder unzureichend überprüfen. Statistiker betonen dagegen bei der Textanalyse in der Regel den Aspekt der Anwendung statistischer Methoden und verzichten meist auf eine tiefergehende Interpretation der gewonnenen Ergebnisse. Außerdem sind viele der Hypothesen, die von Statistikern formuliert werden, vom Standpunkt der Linguistik oder der Literaturwissenschaft aus relativ trivial (vgl. PAUL 1976:68).

Die angedeutete Schwierigkeit dürfte, wie auch die Entwicklung in anderen empirischen Wissenschaften zeigt, nur dadurch zu überwinden sein, daß die Textwissenschaftler sich gründlicher als bisher mit der Statistik beschäftigen und daß sie darüber hinaus bei der Lösung ihrer Probleme so weit wie möglich mit Statistikern zusammenarbeiten.

11.4. Skalierung und Messung textueller Eigenschaften

Im Zusammenhang mit der Verwendung statistischer Modelle bei der Textanalyse spielt das Problem der Skalierung textueller Eigenschaften und Relationen eine bedeutende Rolle.¹⁶

Der Begriff der Skalierung kann folgendermaßen definiert werden (vgl. hierzu KLIEMANN/MÜLLER 1973:172ff; LEINFELLNER 1967:131ff):

Sei M_0 eine Menge empirischer Objekte mit den Relationen $R_1, R_2, \dots, R_n$. Dann heißt

$$\mathcal{M}_0 = \langle M_0, R_1, R_2, \dots, R_n \rangle$$

ein empirisches relationales System (TARSKI) oder auch ein empirisches Relativ. Sei

$$\mathcal{M}_Z = \langle M_Z, S_1, S_2, \dots, S_n \rangle$$

ein numerisches Relativ vom gleichen Typ. Sei

$$\varphi : M_0 \rightarrow M_Z$$

eine strukturerhaltende Abbildung, d.h. ein Morphismus. Unter einer Skalierung soll nun eine isomorphe oder homomorphe Abbildung φ eines empirischen Relativs $\mathcal{M}_0$ in ein numerisches Relativ $\mathcal{M}_Z$ verstanden werden. Die Abbildung φ soll Skalierungsfunktion heißen. Das Ergebnis einer Skalierung ist eine Skala, die durch das geordnete Tripel

$$(\mathcal{M}_0, \mathcal{M}_Z, \varphi)$$

definiert ist.

Die Isomorphie bzw. Homomorphie von $\mathcal{M}_0$ und $\mathcal{M}_Z$ ist eine grundlegende Bedingung für die Anwendung statistischer Modelle. Denn erst dann, wenn man empirisch nachgewiesen hat, daß eine numerische Repräsentation auch die jeweilige Struktur des Untersuchungsgegenstandes abbildet, hat die Anwendung statistischer Modelle überhaupt einen Sinn.

Es gibt eine Vielzahl verschiedener Skalentypen. Am bekanntesten ist die auf STEVENS (1946) zurückgehende Unterscheidung zwischen Nominalskala (kategorische Skala), Ordinalskala, Intervallskala und Verhältnisskala (Ratioskala, Proportionalskala).¹⁷

Die Nominalskala stellt den elementarsten Skalentyp dar und erlaubt lediglich die Klassifikation von Objekten, d.h. die Zerlegung einer Menge von Objekten in Äquivalenzklassen. Sie ist dann realisiert, wenn Objekten zwecks Unterscheidung Namen (Zahlen) zugeordnet werden. Die Ordinalskala erlaubt neben der Feststellung der Gleichheit bzw. Ungleichheit von Objekten auch eine Aussage darüber, welches von zwei Objekten eine größere Menge bestimmter Eigenschaften besitzt. Die Intervallskala erlaubt darüber hinaus auch die Feststellung der Größe des Unterschieds zwischen zwei Objekten. Die Verhältnisskala entspricht dem höchsten Skalierungsniveau. Sie ist zusätzlich noch durch einen absoluten Nullpunkt charakterisiert. Dadurch

werden auch Aussagen über die Gleichheit von Quotienten ermöglicht.

Das Skalierungsniveau ist aufgrund der Isomorphie- bzw. Homomorphiebedingungen von großer Bedeutung für die Anwendbarkeit eines statistischen Modells. So setzt z.B. die Berechnung eines arithmetischen Mittels eine Intervallskalierung voraus. Sind die Eigenschaften einer Intervallskala empirisch nicht erfüllt, kann die Berechnung jedoch zu inadäquaten Ergebnissen führen. Die Anwendung eines statistischen Modells verlangt somit stets eine vorhergehende Bestimmung des Skalierungsniveaus.

Intervall- und Verhältnisskalen werden auch als metrische Skalen bezeichnet. Eine Skalierung auf Nominal- oder Ordinalniveau soll im folgenden als nicht-metrische Skalierung, eine Skalierung auf metrischem Niveau dagegen als metrische Skalierung oder Metrisierung bezeichnet werden. Unter einer Messung soll dagegen der empirische Prozeß der Gewinnung numerischer Daten verstanden werden (vgl. FRIEDRICH/HENNIG 1975:283, STEGMÜLLER 1974:46). Das Resultat einer Messung ist eine Zahl, während das Resultat einer Skalierung (Metrisierung) eine Skala ist. Eine Messung setzt somit stets eine Skalierung voraus.¹⁸

Da die Skalierung nichts anderes als eine spezifische Form der Begriffsbildung ist,¹⁹ können analog zu den verschiedenen Skalentypen auch verschiedene Begriffstypen unterschieden werden.

Eine Nominalskala entspricht einem klassifikatorischen oder qualitativen Begriff, eine Ordinalskala einem komparativen oder topologischen Begriff und eine metrische Skala einem metrischen Begriff (vgl. z.B. OPP 1976:52ff). Metrische und komparative Begriffe bilden zusammen die Klasse der quantitativen Begriffe.²⁰

In der nicht-statistisch orientierten Textwissenschaft (Linguistik/Literaturwissenschaft) wird vorwiegend auf der Nominalskala gemessen, d.h. es werden vor allem qualitative Begriffe verwendet. Wir haben jedoch gesehen, daß umso mehr Eigenschaften und Relationen durch eine Skala abgebildet werden können, je höher das Skalen-

niveau ist. Eine Quantifizierung ist somit stets mit einem Informationsgewinn verbunden, während die Reduktion eines quantitativen Begriffs auf einen qualitativen Begriff stets zu einem Informationsverlust führt. Dies bedeutet, daß man mit quantitativen Begriffen nicht nur alles erreicht, was man mit qualitativen Begriffen erreicht, sondern darüber hinaus noch vieles mehr (vgl. ESSLER 1971:65). Deshalb werden auch Disziplinen wie Linguistik, Literaturwissenschaft, Textwissenschaft und Metrik nicht umhin können, sich verstärkt um die Quantifizierung (Metrisierung) ihres Gegenstandsbereiches zu bemühen.²¹

12. QUANTITATIVE ANALYSE METRISCHER TEXTE

In diesem Kapitel sollen anhand des Gedichts "Erlkönig" von GOETHE und einiger Vergleichstexte einige Methoden der quantitativen Analyse metrischer Texte exemplarisch aufgezeigt werden. Dabei soll gleichzeitig gezeigt werden, daß auch im Fall eines einzelnen Textes mit relativ geringem Umfang die Verwendung quantitativer Methoden durchaus sinnvoll ist und zu tiefergehenden Einsichten in den Untersuchungsgegenstand führt. Ich werde bei der Analyse soweit wie möglich die den verwendeten statistischen Modellen zugrundeliegenden mathematisch-statistischen Überlegungen verdeutlichen. Von einer Darstellung der oft umfangreichen Zwischenrechnungen soll allerdings abgesehen werden.

12.1. Der "Erlkönig" von GOETHE als Analyseobjekt

GOETHEs naturmagische Ballade "Erlkönig" ist 1782, angeregt durch HERDERs Übersetzung der dänischen Volksballade "Erlkönigs Tochter" entstanden. Der "Erlkönig" ist ebenso wie z.B. die Ballade "Der König von Thule" in sog. Volksliedversen abgefaßt (vgl. HEUSLER III 1956:367)¹. Eines der metrisch-formalen Hauptkennzeichen dieser Versart ist die relativ variable Versfüllung und die damit verbundene variable Silbenzahl metrisch gleich gebauter Verse (vgl. PAUL/GLIER 1964:109f). Diese Variabilität kann, wie das Beispiel des "Erlkönig" zeigt, für stilistische Zwecke genutzt werden.²

Der "Erlkönig" - der Text befindet sich im Anhang dieser Arbeit - besteht aus 8 vierzeiligen Strophen mit den Reimpaaren aabb. Das Metrum ist isoakzentuell mit jeweils vier markierten (betonten) Silben pro Vers, wobei die Markierung der letzten Silbe obligatorisch ist. Geht man von einer 'normalsprachlichen' Akzentuierung aus, erhält man folgende, als Grundlage für die quantitative Analyse gedachte Skandierung:

I	1	UMUMUMUM	V	17	UMUMUMUM
	2	UMUMUMUM		18	UMUMUMUM
	3	UMUMUMUM		19	UMUMUMUMUM
	4	UMUMUMUM		20	UMUMUMUMUM
II	5	UMUMUMUMUM	VI	21	UMUMUMUMUM
	6	UMUMUMUM		22	UMUMUMUM
	7	UMUMUMUM		23	UMUMUMUM
	8	UMUMUMUM		24	UMUMUMUMUM
III	9	UMUMUMUM	VII	25	UMUMUMUMUMUM
	10	UMUMUMUM		26	UMUMUMUMUM
	11	UMUMUMUM		27	UMUMUMUMUM
	12	UMUMUMUMUM		28	UMUMUMUM
IV	13	UMUMUMUMUM	VIII	29	UMUMUMUMUM
	14	UMUMUMUMUM		30	UMUMUMUMUM
	15	UMUMUMUM		31	UMUMUMUM
	16	UMUMUMUM		32	UMUMUMUM

Abb. 12.1: Notation des "Erlkönig"

M markiert (Hebung)
U unmarkiert (Senkung)

Ohne das Gedicht weitgehend zu interpretieren, lassen sich in Bezug auf den Inhalt sofort einige relativ eindeutige Feststellungen treffen, z.B. daß die Spannung und Bewegung von der ersten Strophe an wächst, in der siebten Strophe ihren Höhepunkt erreicht und in der achten Strophe ausklingt (vgl. Kap. 12.3). Auch die verschiedenen Personen des Gedichts (Erlkönig, Knabe, Vater, Erzähler) sind inhaltlich unterschiedlich charakterisiert.

Wie wir im Verlauf der quantitativen Analyse sehen werden, können diese inhaltlichen Merkmale zu formalen Stilcharakteristiken in Beziehung gesetzt werden.

12.2. Analyse der Senkungen

12.2.1. Bereits LOTZ (1956) hat im Rahmen seines Notationsvorschlages für germanische Verse - allerdings ohne eine tiefergehende quantitative Analyse vorzunehmen - darauf hingewiesen, daß im "Erlkönig" in bezug auf den Verlauf der Senkungen eine Tendenz zur Strukturierung zu beobachten ist.

Versteht man Strukturiertheit als Nicht-Zufälligkeit der Anordnung - eine Auffassung, die vor allem von osteuropäischen Metrikern vertreten wird - dann kann allein mit Hilfe statistischer Verfahren ermittelt werden, ob ein bestimmtes Textphänomen als zufällig anzusehen ist oder ob eine Tendenz zur nicht-zufälligen Anordnung und damit zur Strukturierung vorliegt.

Untersucht man die Häufigkeit der verschiedenen Senkungstypen vor jeder Hebung, ergeben sich folgende Werte:

Tab. 12.1: Häufigkeit der verschiedenen Senkungstypen von der i-ten Hebung im "Erlkönig"

Senkung j	i-te Hebung				
	i=1	i=2	i=3	i=4	
0	2	0	0	0	2
1	27	22	13	11	73
2	3	9	19	21	52
3	0	1	0	0	1
	32	32	32	32	128

- j = 0 fehlende Senkung
- j = 1 monosyllabische Senkung
- j = 2 bisyllabische Senkung
- j = 3 trisyllabische Senkung

Ein Blick in die Tabelle zeigt, daß erstens der weitaus größte Teil der Senkungen mono- oder bisyllabisch ist, daß zweitens die Häufigkeit der verschiedenen Senkungstypen von der jeweiligen Position im Vers abzuhängen scheint und daß drittens die Anzahl der monosyllabischen Senkungen im Verlauf der Verszeile abnimmt, während die Anzahl der bisyllabischen Senkungen zunimmt.

Es soll nun zuerst mit Hilfe eines χ^2 -Tests geprüft werden, ob ein statistisch signifikanter Zusammenhang zwischen der jeweiligen Position einer Hebung - im folgenden Hebungstyp genannt - und dem jeweiligen Senkungstyp besteht. Um die Voraussetzungen des χ^2 -Tests zu erfüllen - die Erwartungswerte dürfen nicht zu klein sein - legen wir die Häufigkeiten der Senkungen mit 0 und mit 3 Silben mit den Häufigkeiten der mono- bzw. bisyllabischen Senkungen zusammen und unterscheiden nur noch zwischen 'kurzen' und 'langen' Senkungen. Hierbei muß allerdings berücksichtigt werden, daß zum einen j = 0 und j = 1 qualitativ sehr unterschiedlich sind und daß zum anderen gerade seltene Ereignisse, die unter statistischen Gesichtspunkten als zufällig zu betrachten sind, häufig die durch den Kontext geschaffene Erwartungshaltung des Rezipienten durchbrechen und dadurch eine ästhetische Wirkung ausüben können (vgl. RIFFATERRE 1960, LEVIN 1963, 1965).

Um die Hypothese eines Zusammenhangs zwischen Hebungstyp und Senkungstyp prüfen zu können, muß diese Hypothese zuerst in die Sprache des statistischen Modells übersetzt werden (vgl. Kap. 11.3). Es ist in der Teststatistik üblich, zwischen der sog. Nullhypothese H_0 und der sog. Alternativhypothese H_1 zu unterscheiden.³ Die Alternativhypothese entspricht inhaltlich in der Regel der zu prüfenden Forschungshypothese. Getestet wird jeweils die Nullhypothese gegen die Alternativhypothese. Spricht das Testresultat gegen die Nullhypothese wird die Nullhypothese verworfen und die Alternativhypothese akzeptiert. Spricht das Testresultat für die Nullhypothese wird die Alternativhypothese verworfen und die Nullhypothese akzeptiert. Eine Entscheidung für eine Hypothese bedeutet jedoch nicht, daß damit etwa die Gültigkeit der entsprechenden Hypothese ein für alle mal bewiesen ist. Jede Entscheidung über Annahme oder Verwerfung einer

Hypothese ist mit einer Irrtumswahrscheinlichkeit verbunden. Diese kann jedoch, und dies ist ein wichtiger Unterschied zu Entscheidungen, die ohne Verwendung von statistischen Methoden getroffen werden, genau bestimmt werden (vgl. Kap. 11.3).

Die Null- und die Alternativhypothese könnten im vorliegenden Fall folgendermaßen formuliert werden:

H_0 : Es besteht kein Zusammenhang zwischen Hebungstyp und Senkungstyp.

H_1 : Es besteht ein Zusammenhang zwischen Hebungstyp und Senkungstyp.

Unter der Nullhypothese ist zu erwarten, daß die Wahrscheinlichkeiten für das Auftreten der Ereignisse 'Senkung (S)' und 'Hebung (H)' voneinander unabhängig sind. Dies bedeutet, daß die Wahrscheinlichkeit für das gemeinsame Auftreten der Ereignisse H_i und S_j gleich dem Produkt der Einzelwahrscheinlichkeiten sein muß. In der Sprache des statistischen Modells lauten H_0 und H_1 jetzt folgendermaßen:

$$H_0: P(H_i S_j) = P(H_i)P(S_j)$$

$$H_1: \neg (P(H_i S_j) = P(H_i)P(S_j))$$

Mit Hilfe des χ^2 -Tests kann geprüft werden, ob Abweichungen zwischen den beobachteten Häufigkeiten und den unter der Nullhypothese erwarteten Häufigkeiten dem Zufall zuzuschreiben sind oder ob ein systematischer Einfluß vorliegt (zum χ^2 -Test vgl. z.B. HAYS 1973:717ff).

Bezeichnen wir mit O_{ij} die beobachteten Häufigkeiten mit E_{ij} die erwarteten Häufigkeiten, lautet die Prüfstatistik:

$$\chi^2 = \sum_{i=1}^k \sum_{j=1}^r \frac{(O_{ij} - E_{ij})^2}{E_{ij}} \quad (12.1)$$

Berücksichtigt man, daß

$$\sum_{i=1}^k \sum_{j=1}^r O_{ij} = \sum_{i=1}^k \sum_{j=1}^r E_{ij} = n$$

ist, erhält man aus (12.1) die rechnerisch etwas einfachere Formel (12.2):

$$\chi^2 = \sum_{i=1}^k \sum_{j=1}^r \frac{O_{ij}^2}{E_{ij}} - n \quad (12.2)$$

Für die erwarteten Häufigkeiten gilt unter der Nullhypothese

$$E_{ij} = P(H_i)P(S_j) \cdot n$$

Dementsprechend lauten die Erwartungswerte für die kurzen Senkungen

$$E_{i1} = \frac{75}{128} \cdot \frac{32}{128} \cdot 128 = 18,75$$

und für die langen Senkungen

$$E_{i2} = \frac{53}{128} \cdot \frac{32}{128} \cdot 128 = 13,25.$$

Die erwarteten Häufigkeiten sind zusammen mit den beobachteten Häufigkeiten in der folgenden Tabelle aufgeführt:

Tab. 12.2: Beobachtete und erwartete Häufigkeiten für kurze und lange Senkungen im "Erlikönig"

Senkung	i-te Hebung							
	O_{1j}	E_{1j}	O_{2j}	E_{2j}	O_{3j}	E_{3j}	O_{4j}	E_{4j}
kurz	29	18,75	22	18,75	13	18,75	11	18,75
lang	3	13,25	10	13,25	19	18,75	21	13,25
	32		32		32		32	128

Wie man sieht, weichen die beobachteten Häufigkeiten vor allem vor der ersten Hebung und etwas weniger deutlich vor der vierten und vor der dritten Hebung von den Erwartungswerten ab. Vor der zweiten Hebung besteht dagegen eine relativ große Übereinstimmung zwischen beobachteten und erwarteten Häufigkeiten. Der χ^2 -Test ergibt einen Wert von $\chi^2 = 26,89$. Wie man zeigen kann, folgt die Stichprobenverteilung

von χ^2 für $n \rightarrow \infty$ (d.h. für große Erwartungswerte) einer χ^2 -Verteilung mit $(k - 1) (r - 1)$ Freiheitsgraden (FG). Da für FG = $(4 - 1) (2 - 1) = 3$ $P(\chi^2 \geq 26,89) < 0,0001$ ist, kann mit einer Irrtumswahrscheinlichkeit von weniger als 0,01 % die Nullhypothese, d.h. die Hypothese der Unabhängigkeit zwischen Hebungstyp und Senkungstyp verworfen werden.

12.2.2. Der Zusammenhang zwischen Hebungstyp und Senkungstyp kann über eine Regressionsanalyse genauer erfaßt werden. Die allgemeine Form einer Regressionsgleichung für die Regression von y auf x lautet im Fall eines linearen Modells

$$\hat{y} = a_{yx} + b_{yx}x \quad (12.3)$$

Es gibt eine Reihe von Methoden, mit denen die Parameter einer solchen Regressionsgleichung geschätzt werden können. Bei der Methode der kleinsten Quadrate geht man davon aus, daß die quadrierten Abweichungen der geschätzten y -Werte von den beobachteten y -Werten ein Minimum bilden sollen, d.h. daß die Bedingung

$$\sum_{i=1}^n (y_i - \hat{y}_i) = \text{Min} \quad (12.4)$$

gelten soll.

Ersetzen wir in (12.4) $\hat{y}_i$ durch $a + bx_i$, erhalten wir

$$q = f(a, b) = \sum_{i=1}^n (y_i - a - bx_i)^2.$$

Damit diese Funktion ein Minimum hat, muß

$$\frac{\partial q}{\partial b} = 0 \quad \text{und} \quad \frac{\partial q}{\partial a} = 0$$

sein. Außerdem müssen noch folgende Bedingungen erfüllt sein:

$$\frac{\partial^2 q}{\partial a^2} \frac{\partial^2 q}{\partial b^2} > \left(\frac{\partial^2 q}{\partial a \partial b} \right)^2 \quad \text{und} \quad \frac{\partial^2 q}{\partial a^2} > 0$$

Wir erhalten die partiellen Ableitungen (der Summationsbereich ist im folgenden 1 bis n)

$$\frac{\partial q}{\partial b} = 2 \sum (y_i - a - bx_i)(-x_i) = 0$$

$$\sum x_i y_i = a \sum x_i + b \sum x_i^2 \quad (12.5)$$

$$\frac{\partial q}{\partial a} = -2 \sum (y_i - a - bx_i) = 0$$

$$\sum y_i = na + b \sum x_i. \quad (12.6)$$

Da auch die weiteren Bedingungen erfüllt sind, sind (12.5) und (12.6) tatsächlich Minima.

Das aus den sog. Normalgleichungen (12.5) und (12.6) gebildete Gleichungssystem hat die folgenden Lösungen:

$$b = \frac{n \sum x_i y_i - \sum x_i \sum y_i}{n \sum x_i^2 - (\sum x_i)^2} \quad (12.7)$$

(12.7) kann auch als

$$b = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} = \frac{s_{xy}}{s_x^2} \quad (12.8)$$

geschrieben werden.

Für a erhalten wir

$$a = \frac{\sum x_i^2 \sum y_i - \sum x_i \sum x_i y_i}{n \sum x_i^2 - (\sum x_i)^2} \quad (12.9)$$

Dividiert man (12.6) durch n , ergeben sich für a folgende weitere Formeln

$$a = \bar{y} - b\bar{x} \quad (12.10)$$

und

$$a = \frac{\sum y_i - b \sum x_i}{n} \quad (12.11)$$

Liegt eine $k \times 2$ -Felder-Tafel vor (vgl. Tab. 12.2) und bilden die k Merkmale eine 'natürliche' Reihenfolge (z.B. erster, zweiter, ..., n -ter Hebungstyp), dann kann man die Reihenfolge der k Merkmale durch entsprechende Punktwerte ausdrücken und die Regression der beobachteten relativen Häufigkeiten auf die durch die Punktwerte gewichteten Merkmale berechnen. Bezeichnet man die Punktwerte mit x , dann lautet der Regressionskoeffizient für die gewichtete Regression der relativen Häufigkeiten $\hat{p}$ auf die Punktwerte x

$$b = \frac{\sum n_i (\hat{p}_i - \bar{p}) (x_i - \bar{x})}{\sum n_i (x_i - \bar{x})^2} \quad (12.12)$$

bzw.

$$b = \frac{\sum x_i y_i - (\sum n_i x_i \sum y_i) / n}{\sum n_i x_i^2 - (\sum n_i x_i)^2 / n}, \quad (12.13)$$

wobei $\bar{p} = \frac{\sum y_i}{n}$, $\hat{p}_i = \frac{y_i}{n_i}$, $\bar{x} = \frac{\sum n_i x_i}{n}$ und $n = \sum n_i$ ist.

Für den Achsenabschnitt erhalten wir

$$a = \frac{\sum y_i - b \sum n_i x_i}{n}. \quad (12.14)$$

Mit Hilfe der abgeleiteten Formeln können wir nun die Regression der langen Senkungen p_i auf den Hebungstyp x_i berechnen. Die dazu benötigten Werte finden sich in der Tabelle 12.3. Anstelle der Zahlenwerte 1,2,3,4 hätten wir zur Bezeichnung der Reihenfolge der Hebungstypen z.B. auch die Werte -1,0,1,2 benutzen können. Diese Gewichtung hätte die folgenden Rechnungen etwas vereinfacht; die Ergebnisse wären jedoch die gleichen geblieben.

Setzen wir die entsprechenden Zahlen in die Formeln (12.13) und (12.14) ein, erhalten wir den Regressionskoeffizienten

$$b = \frac{164 - 320(53)/128}{960 - 320^2/128} = 0,1969$$

und den Achsenabschnitt

$$a = \frac{53 - 0,1969(320)}{128} = -0,0781.$$

Tab. 12.3: Berechnung der Regression der langen Senkungen $\hat{p}_i$ auf den Hebungstyp x_i

x_i	y_i	n_i	$\hat{p}_i$	$x_i y_i$	$n_i x_i$	$n_i x_i^2$
1	3	32	0,0938	3	32	32
2	10	32	0,3125	20	64	128
3	19	32	0,5938	57	96	288
4	21	32	0,6563	84	128	512
	53	128		164	320	960

Die Regressionsgerade lautet demzufolge

$$p' = -0,0781 + 0,1969x \quad (12.15)$$

Da zwischen den langen Senkungen ($\hat{p}$) und den kurzen Senkungen ($\hat{q}$) die funktionale Beziehung $\hat{q} = 1 - \hat{p}$ besteht, ergibt sich aus (12.15) für die Regression der kurzen Senkungen auf den Hebungstyp die Gerade

$$q' = 1,0781 - 0,1969x. \quad (12.16)$$

Setzen wir die Werte für x in die beiden Regressionsgeraden ein, erhalten wir folgende Schätzwerte:

Tab. 12.4: Beobachtete und geschätzte Werte für die relativen Häufigkeiten von langen und kurzen Senkungen pro Vers im "Erk König"

x_i	$\hat{p}_i$	p'_i	$\hat{q}_i$	q'_i
1	0,0938	0,1188	0,9062	0,8812
2	0,3125	0,3156	0,6875	0,6844
3	0,5938	0,5125	0,4062	0,4875
4	0,6563	0,7094	0,3437	0,2906

Die Daten der Tabelle 12.4 sind in der folgenden Abbildung dargestellt:

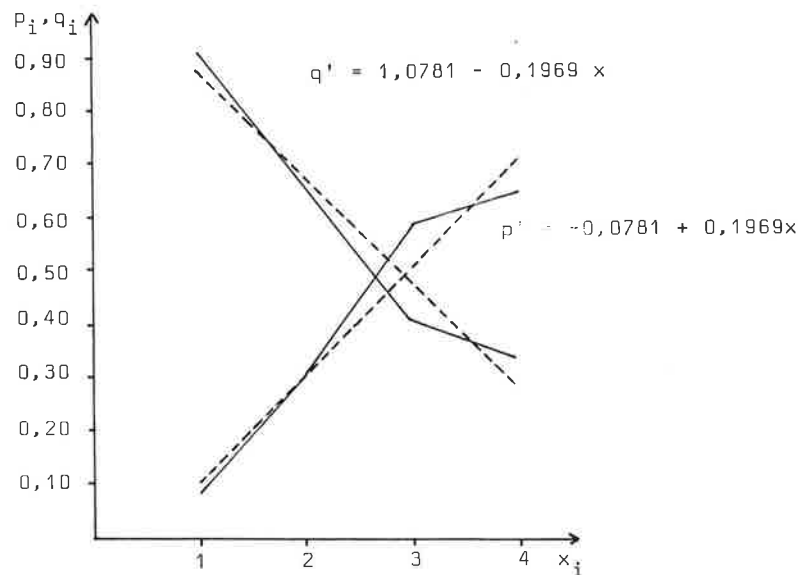


Abb. 12.2: Verlauf der Senkungen im Vers

— empirische Kurven
 ---- Regressionsgeraden

Wie aus der Tabelle 12.4 und aus der Abbildung 12.2 ersichtlich ist, beschreiben die Regressionsgeraden (12.15) und (12.16) ziemlich genau den Verlauf von langen bzw. kurzen Senkungen im Vers. Die Regressionskoeffizienten zeigen dabei, daß der Anteil der langen Senkungen vor jeder Hebung um durchschnittlich 19,69 % zunimmt, während der Anteil der kurzen Senkungen um durchschnittlich 19,69 % abnimmt.

12.2.3. Wir wollen jetzt testen, ob diese Tendenz signifikant ist oder lediglich dem Zufall zuzuschreiben ist. Dazu zerlegen wir das $\chi^2 = 26,89$, das wir bei der Prüfung des Zusammenhangs zwischen Hebungs- und Senkungstyp erhalten haben, in zwei Anteile: Der eine Anteil entspricht den linear ansteigenden relativen Häufigkeiten, der restliche Anteil entfällt auf die Unterschiede zwischen den beobachteten und den linear ansteigend vorausgesetzten theoretischen Häufigkeiten. Dieses von COCHRAN (1954) vorgeschlagene Verfahren erlaubt, den Anteil der linearen Regression an der Gesamtvariation von dem Anteil abzugrenzen, der durch die Abweichungen von der Regressionsgeraden bestimmt ist (vgl. auch ARMITAGE 1955).⁴

Das χ^2 für die lineare Regression erhalten wir, indem wir in der Formel für den gewichteten Regressionskoeffizienten (vgl. Formel 12.13) den Zähler quadrieren und durch $\bar{p}(1-\bar{p})$ dividieren:

$$\chi^2_{\text{Regr.}} = \frac{[\sum x_i y_i - (\sum n_i x_i \sum y_i) / r]^2}{\bar{p}(1-\bar{p}) [\sum n_i x_i^2 - (\sum n_i x_i)^2 / n]} \quad (12.17)$$

Die Ausdrücke in eckigen Klammern haben wir bereits bei der Berechnung des Regressionskoeffizienten ermittelt. Das χ^2 für die Regression lautet somit:

$$\chi^2_{\text{Regr.}} = \frac{31,5^2}{0,4141(0,5859)160} = 25,5614$$

Für $FG = 1$ ist $P(\chi^2 \geq 25,56) \approx 0,0000004$. Die Regression ist damit hochsignifikant. Insgesamt erhalten wir folgende χ^2 -Anteile:

Tab. 12.5: Anteil der linearen Regression und der Abweichungen von der Regression am Gesamt- χ^2 für die Daten der Tabelle 12.3

Variationsursache	χ^2	FG	P
Lineare Regression	25,56	1	$0,4 \cdot 10^{-6}$
Abweichungen von der Regression	1,33	2	0,2401
Insgesamt	26,89	3	$< 0,0001$

Die Tabelle 12.5 zeigt den entscheidenden Anteil der linearen Regression an der Gesamtvariation. Die Abweichungen von der Regression können dagegen als zufällig betrachtet werden. Die Anzahl der langen Senkungen wächst somit regelmäßig im Verlauf des Verses, während die Anzahl der kurzen Senkungen entsprechend abnimmt. Wie der χ^2 -Test zeigt, wird diese Tendenz hinlänglich gut durch die Regressionsgeraden (12.15) und (12.16) beschrieben.

12.2.4. Es soll jetzt geprüft werden, ob in der Ballade "Der Totentanz", die ebenfalls von GOETHE stammt, eine ähnliche Tendenz vorliegt.

Die Ballade "Der Totentanz" ist im Jahre 1813 entstanden, also 31 Jahre später als der "Erlkönig". Sie besteht aus 7 Strophen zu je 7 Versen. 4 Verse einer Strophe weisen jeweils 4 Hebungen, die übrigen 3 Hebungen auf. Ich beschränke mich beim Vergleich auf die Verse mit 4 Hebungen. Die entsprechenden Daten finden sich in der Tabelle 12.6.

Tab. 12.6: Häufigkeit der verschiedenen Senkungstypen vor der i-ten Hebung in "Der Totentanz" von GOETHE

Senkung j	i-te Hebung			
	i=1	i=2	i=3	i=4
1	28	0	2	0
2	0	28	26	28

Wie die Tabelle zeigt, weisen die 4-hebigen Verse dieser Ballade im Gegensatz zum "Erlkönig" vor den einzelnen Hebungen fast keine metrische Variation auf. Vor der ersten Hebung werden ausschließlich monosyllabische Senkungen, vor den übrigen Hebungen dagegen fast ausschließlich bisyllabische Senkungen verwendet. Der Unterschied zum "Erlkönig" ist so deutlich ausgeprägt, daß auf einen Signifikanztest verzichtet werden kann. Die Verteilung der Senkungen pro Vers ist somit zumindest in bezug auf die beiden untersuchten Balladen ein geeigneter Stildiskriminator (vgl. Kap. 11.2).

"Der Totentanz" und der "Erlkönig" unterscheiden sich jedoch nicht nur aufgrund der Verteilung der Senkungen im Vers, sondern auch aufgrund der relativen Häufigkeiten von kurzen und langen Senkungen im Gesamttext. Berücksichtigen wir auch die 3-hebigen Verse in "Der Totentanz" erhalten wir folgende Daten:

Tab. 12.7: Absolute (f_i) und relative (%) Häufigkeiten von kurzen und langen Senkungen im "Erlkönig" und in "Der Totentanz"

	Erlkönig		Totentanz		
	f_i	%	f_i	%	
kurz	75	58,59	72	36,73	147
lang	53	41,41	124	63,27	177
	128		196		324

Wir testen zuerst, ob das Verhältnis von kurzen und langen Senkungen in beiden Gedichten einer Gleichverteilung folgt, d.h. ob $H_0: p = 0,5$ gilt. Da keine begründete Annahme über die Richtung des jeweiligen Unterschieds gemacht werden kann, ist ein zweiseitiger Test angebracht, d.h. es wird H_0 gegen die Alternativhypothese $H_1: p \neq 0,5$ getestet. Für $npq \geq 9$ ist ein Test mit Hilfe der Normal-

verteilung möglich. Die Teststatistik lautet in allgemeiner Form

$$\hat{z} = \frac{X - E(X)}{\sqrt{V(X)}} \quad (12.18)$$

Diese Statistik folgt für $n \rightarrow \infty$ einer Normalverteilung mit dem Mittelwert $\mu' = 0$ und der Standardabweichung $\sigma = 1$. Da es sich im vorliegenden Fall um binomialverteilte relative Häufigkeiten mit $E(\hat{p}) = p$, $V(\hat{p}) = \frac{pq}{n}$ handelt, gilt

$$\hat{z} = \frac{\hat{p} - p}{\sqrt{\frac{pq}{n}}} = \frac{(\hat{p} - p) \sqrt{n}}{\sqrt{pq}} \quad (12.19)$$

Wir erhalten folgende Werte:

"Erlkönig":

$$\hat{z} = \frac{(0,5859 - 0,5) \sqrt{128}}{\sqrt{0,5(0,5)}} = 1,9445$$

"Der Totentanz":

$$\hat{z} = \frac{(0,3637 - 0,5) \sqrt{196}}{\sqrt{0,5(0,5)}} = 3,7156$$

Die beiden Werte besagen, daß im "Erlkönig" eine nahezu signifikante Tendenz ($P = 0,052$) besteht, insgesamt gesehen mehr kurze als lange Senkungen zu verwenden. "Der Totentanz" weist dagegen eine hochsignifikante Tendenz ($P = 0,0002$) auf, vor allem lange Senkungen zu verwenden.

12.2.5.1. Wir wollen jetzt prüfen, wie stark der Zusammenhang zwischen den Variablen 'Gedicht' und 'Senkungstyp' ist. Dazu testen wir den Zusammenhang zuerst mit Hilfe eines χ^2 -Tests. Bezeichnet man in einer 4-Feldertafel die Häufigkeiten der Hauptdiagonalen mit a und d, die Häufigkeiten der Nebendiagonalen mit b und c, die Randsummen mit a+b, a+c, b+d, c+d und a+b+c+d = n, erhält man nach einigen einfachen, jedoch langwierigen Umformungen aus der üblichen χ^2 -Formel (vgl. Formel 12.1) die folgende χ^2 -Formel,

die speziell für 4-Felder-Tafeln gilt und rechnerisch wesentlich einfacher ist:

$$\chi^2 = \frac{n(ad - bc)^2}{(a+b)(a+c)(b+d)(c+d)} \quad (12.20)$$

Für die Tabelle 12.7 lautet das χ^2 :

$$\chi^2 = \frac{324(75 \cdot 124 - 72 \cdot 53)^2}{147(128)196(177)} = 14,9274$$

Wie aufgrund der Ergebnisse der Einzeltests zu erwarten war, besteht in bezug auf die Anzahl der kurzen und langen Senkungstypen ein hochsignifikanter Unterschied zwischen den beiden Gedichten ($P \ll 0,01$).

Mit dieser Feststellung wird jedoch nichts über die Stärke des Zusammenhangs zwischen den Untersuchungsvariablen 'Senkungstyp' und 'Gedicht' ausgesagt. Denn wie die χ^2 -Formel (12.2) zeigt, hängt der numerische Wert der χ^2 -Statistik sowohl von der Größe der Differenz zwischen beobachteten und erwarteten Häufigkeiten als auch von der Stichprobengröße ab. Außerdem spielt auch die Anzahl der Felder der jeweiligen Kontingenztafel eine Rolle. Dies bedeutet, daß mit wachsendem Stichprobenumfang auch bei einem sehr geringen Zusammenhang zwischen den Untersuchungsvariablen ein hochsignifikanter Wert erreicht werden kann. Außerdem ist auch ein direkter Vergleich von χ^2 -Werten, die auf einem unterschiedlichen Stichprobenumfang beruhen, nicht möglich.

Diesen beiden Tatsachen, die bei der Interpretation des χ^2 -Tests häufig nicht berücksichtigt werden, kann dadurch Rechnung getragen werden, daß wir das χ^2 durch sein jeweiliges Maximum dividieren. Auf diese Weise erhalten wir einen Assoziationskoeffizienten, der im Intervall (0,1) variiert. Der Koeffizient nimmt den Wert 0 an, wenn beobachtete und erwartete Häufigkeiten gleich sind, d.h. wenn zwischen den Untersuchungsvariablen absolute Unabhängigkeit besteht. Der Wert 1 wird dann erreicht, wenn das χ^2 seinen jeweiligen Maximalwert annimmt. Dies ist z.B. der Fall, wenn in einer symmetrischen $k \times k$ -Feldertafel lediglich die Felder einer einzigen Diagonalen besetzt sind.

In diesem Fall herrscht eine absolute Abhängigkeit zwischen den Untersuchungsvariablen.

Wie man leicht zeigen kann, besteht zwischen der Größe einer $k \times r$ -Feldertafel, dem Stichprobenumfang n und dem maximalen χ^2 -Wert folgender Zusammenhang:

$$\max \chi^2 = \min (k - 1, r - 1) \cdot n.$$

Dies bedeutet, daß z.B. im Fall einer 4×3 -Feldertafel das χ^2 höchstens einen Wert von $(r - 1)n = 2n$ erreichen kann. Wir erhalten folgenden Assoziationskoeffizienten:

$$C = \frac{\chi^2}{\max \chi^2} = \frac{\chi^2}{n \cdot \min (k - 1, r - 1)}$$

Die Wurzel aus C entspricht dem sog. CRAMÉRSchen Kontingenzkoeffizienten C_c :

$$C_c = \sqrt{\frac{\chi^2}{n \cdot \min (k - 1, r - 1)}} \quad (12.21)$$

Dieser Koeffizient, der auch direkt aus dem PEARSON-BRAVAISSchen Produkt-Moment-Korrelationskoeffizienten hergeleitet werden kann, geht im Fall einer 2×2 -Feldertafel in den Kontingenzkoeffizienten ϕ über:

$$\phi = \sqrt{\frac{\chi^2}{n}} \quad (12.22)$$

Setzt man in dieser Formel für das χ^2 die χ^2 -Formel für 2×2 -Feldertafeln (vgl. 12.20) ein, erhält man

$$\phi = \frac{ad - bc}{\sqrt{(a + b)(a + c)(c + d)(b + d)}} \quad (12.23)$$

Für den Zusammenhang zwischen den Variablen 'Gedicht' und 'Senkungstyp' hatten wir ein $\chi^2 = 14,9274$ ermittelt. Dem entspricht ein

$$\phi = \sqrt{\frac{14,9274}{324}} = 0,2146.$$

Dieser Wert ist trotz des hochsignifikanten χ^2 -Wertes relativ gering. Es liegt somit lediglich ein ziemlich schwacher Einfluß der Variablen 'Gedicht' auf die Variable 'Senkungstyp' vor.

12.2.5.2. Der Phi-Koeffizient ist ein relativ häufig verwendetes Maß für den Zusammenhang zwischen dichotomen nominalskalierten Variablen. Der Phi-Koeffizient und der CRAMÉRSche Kontingenzkoeffizient haben vor allem den Vorteil, daß sie relativ leicht über die χ^2 -Statistik berechnet und auf Signifikanz getestet werden können. Sie haben allerdings den entscheidenden Nachteil, daß sie keine eindeutige inhaltliche Interpretation erlauben (zur Kritik an den χ^2 -basierten Assoziationsmaßen vgl. vor allem GOODMAN/KRUSKAL 1954, COSTNER 1965). Eine klare inhaltliche Interpretation ist dagegen bei den sog. PRE-Maßen ("proportional reduction in error measures") möglich, die in jüngerer Zeit immer mehr an Bedeutung gewinnen. Diese Maße, die vor allem auf GOODMAN/KRUSKAL (1954) zurückgehen, beschreiben den Grad, in dem uns die Kenntnis der einen Variablen die andere Variable vorherzusagen hilft.

Ein PRE-Maß für den Zusammenhang zwischen nominalskalierten Daten ist das "Maß der prädiktiven Assoziation" λ (vgl. zum folgenden HAYS 1973:745ff, BENNINGHAUS 1974:87ff und 125ff). Dieses Maß, das bei jeder Tabellengröße berechnet werden kann und zwischen 0 und 1 variiert, setzt wie alle PRE-Maße

- (1) eine Regel für die Vorhersage der abhängigen Variablen auf der Basis ihrer eigenen Verteilung,
- (2) eine Regel für die Vorhersage der abhängigen Variablen auf der Basis der unabhängigen Variablen und
- (3) eine Fehlerdefinition voraus.

Lambda hat wie alle PRE-Maße die allgemeine Form

$$\text{PRE-Maß} = \frac{E_1 - E_2}{E_1} \quad (12.24)$$

wobei E_1 die Anzahl der Fehler nach Regel 1 und E_2 die Anzahl der Fehler nach Regel 2 bezeichnet.

Für Lambda lauten die entsprechenden Regeln:

- (1) Es wird der Modalwert der Marginalverteilung der abhängigen Variablen bestimmt. Dieser Wert erlaubt die 'beste' Voraussage der abhängigen Variablen auf der Basis ihrer eigenen Verteilung.
- (2) Es werden für jede Kategorie der unabhängigen Variablen die Modalwerte der gemeinsamen Verteilung bestimmt. Diese Werte erlauben die 'beste' Vorhersage der abhängigen Variablen auf der Basis der unabhängigen Variablen.
- (3) Jeder von den Vorhersageregeln abweichende Fall ist ein Fehler.

Lambda ist ein asymmetrisches Maß, d.h. es können für jede Tabelle zwei Lambda-Werte berechnet werden und zwar einmal mit der abhängigen Variablen am Tabellenrand und zum anderen mit der abhängigen Variablen im Tabellenkopf. Außerdem kann noch, wenn keine der beiden Variablen als von der anderen unabhängig angesehen werden kann, aus der Kombination der beiden Werte ein weiteres 'symmetrisches' Lambda berechnet werden. Im folgenden berücksichtige ich lediglich die Möglichkeit, daß die abhängige Variable - wie es meist der Fall ist - die Zeilenvariable ist. Ist die abhängige Variable die Spaltenvariable, verläuft die Argumentation analog.

Die Anzahl der Fehler, die nach Regel 1 gemacht werden können, beträgt für den Fall, daß die abhängige Variable die Zeilenvariable ist,

$$E_1 = n - \max n_{i.} \quad (12.25)$$

wobei $\max n_{i.}$ die Anzahl der Untersuchungseinheiten in der modalen Kategorie der Zeilenvariablen bezeichnet. In der Tabelle 12.7 beträgt der Modalwert der abhängigen Variablen 177. Der Vorhersagefehler nach Regel 1 ist somit

$$E_1 = 324 - 177 = 147.$$

Für die Anzahl der Fehler, die nach Regel 2 gemacht werden können, gilt

$$E_2 = \sum_{j=1}^r (n_{.j} - \max n_{ij}). \quad (12.26)$$

Da in Tabelle 12.7 der Modalwert der Kategorie "Erlkönig" 75 beträgt und der Modalwert der Kategorie "Totentanz" 124 ist, lautet E_2

$$E_2 = 128 - 75 + 196 - 124 = 125.$$

Wir erhalten für λ folgenden Wert:

$$\lambda = \frac{E_1 - E_2}{E_1} = \frac{147 - 125}{147} = 0,1497$$

Dieser Wert besagt, daß bei Kenntnis der Variablen 'Gedicht' der Fehler bei der Vorhersage der Werte für die abhängige Variable 'Senkungstyp' um 14,97 % geringer ist als bei einer Vorhersage, die sich lediglich auf die Werte der abhängigen Variablen stützt. Der Vorteil einer solchen eindeutigen inhaltlichen Interpretation gegenüber der relativ vagen Interpretation des CRAMÉRSchen Kontingenzkoeffizienten bzw. des Phi-Koeffizienten ist offensichtlich.

Es ist auch möglich, Lambda ohne eine vorhergehende Bestimmung von E_1 und E_2 direkt zu berechnen. Die Formel hierfür ergibt sich aus den Formeln (12.24), (12.25) und (12.26):

$$\lambda = \frac{E_1 - E_2}{E_1} = \frac{n - \max n_{i.} - \sum_{j=1}^r (n_{.j} - \max n_{ij})}{n - \max n_{i.}} \quad (12.27)$$

$$\lambda = \frac{\sum_{j=1}^r \max n_{ij} - \max n_{i.}}{n - \max n_{i.}}$$

12.2.6. Wir haben festgestellt, daß im "Erlkönig" insgesamt gesehen vor allem kurze Senkungen verwendet werden. Es soll jetzt untersucht werden, ob diese Tendenz in allen Strophen in gleicher Weise ausgeprägt ist. Die dazu benötigten Daten finden sich in Tabelle 12.8.

Tab. 12.8: Anzahl der kurzen und langen Senkungen pro Strophe im "Erlkönig"

Senkung	Strophe								
	1	2	3	4	5	6	7	8	
kurz	12	12	12	9	7	8	5	10	75
lang	4	4	4	7	9	8	11	6	53
	16	16	16	16	16	16	16	16	128

Es zeigt sich, daß die Anzahl der langen Senkungen bis zur 7. Strophe anwächst und in der 8. Strophe wieder abnimmt. Die 7. Strophe enthält außerdem die einzige trisyllabische Senkung des Gedichts.

Wir überprüfen zuerst den Zusammenhang der Variablen 'Strophe' und 'Senkungstyp' mit Hilfe des üblichen χ^2 -Tests. Im Fall von umfangreicheren $k \times 2$ - bzw. $2 \times r$ -Feldertafeln empfiehlt es sich, zur Rechenerleichterung die χ^2 -Formel von BRANDT/SNEDECOR zu benutzen. Diese Formel erhält man durch folgende Überlegung (vgl. HALD 1967:711ff):

Geht man davon aus, daß jede der k Stichproben einer Normalverteilung mit den Parametern $(0, 1)$ folgt, gilt (vgl. 12.19)

$$\hat{z}_i = \frac{\hat{p}_i - p}{\sqrt{\frac{p(1-p)}{n_i}}}$$

wobei $\hat{p}_i = \frac{x_i}{n_i}$ ist.

Da z^2 einem χ^2 mit $FG = 1$ entspricht und da χ^2 -verteilte Zufallsvariablen additiv sind, gilt für k Stichproben

$$\chi^2 = \sum_{i=1}^k \hat{z}_i^2 = \sum_{i=1}^k \frac{n_i (\hat{p}_i - p)^2}{p(1-p)}$$

Schätzen wir den Parameter p durch

$$\frac{\sum n_i \hat{p}_i}{\sum n_i} = \bar{p}$$

erhalten wir die gesuchte χ^2 -Formel

$$\chi^2 = \frac{1}{p(1-p)} \sum_{i=1}^k n_i (\hat{p}_i - \bar{p})^2 \quad (12.28)$$

mit $FG = k - 1$.

Setzen wir

$$\sum x_i = x \text{ und } \bar{p} = \frac{\sum n_i \hat{p}_i}{\sum n_i} = \frac{\sum x_i}{n} = \frac{x}{n},$$

folgt aus (12.28) die BRANDT/SNEDECOR-Formel für absolute Häufigkeiten:

$$\begin{aligned} \chi^2 &= \frac{1}{\frac{x}{n}(1 - \frac{x}{n})} \sum_{i=1}^k n_i \left(\frac{x_i}{n_i} - \frac{x}{n} \right)^2 \\ &= \frac{n^2}{x(n-x)} \left[\sum_{i=1}^k \frac{x_i^2}{n_i} - \frac{x^2}{n} \right] \quad (12.29) \end{aligned}$$

Für die Tabelle 12.8 ergibt sich nach dieser Formel folgender χ^2 -Wert:

$$\chi^2 = \frac{128^2}{75(53)} \left(\frac{12^2}{16} + \dots + \frac{10^2}{16} - \frac{75^2}{128} \right) = 12,3331$$

Da $P(X_7^2 \geq 12,3331) \approx 0,0921$ ist, muß bei einem Signifikanzniveau von $P = 0,05$ die Nullhypothese, d.h. die Hypothese der Unabhängigkeit der Variablen 'Strophe' und 'Senkungstyp' beibehalten werden.

Bei diesem Ergebnis muß jedoch berücksichtigt werden, daß der Nullhypothese keine spezifische Alternativhypothese gegenübergestellt worden ist. Dieses beim X^2 -Test übliche Vorgehen kann zur Folge haben, daß man ein nichtsignifikantes Resultat erhält, obwohl die Nullhypothese in Wirklichkeit falsch ist (vgl. COCHRAN 1954). Prüft man dagegen die Daten mit Hilfe des COCHRANschen X^2 -Tests (vgl. Formel 12.17) auf lineare Regression, erhält man die folgenden X^2 -Werte:

Tab. 12.9: Anteil der linearen Regression und der Abweichungen von der Regression am Gesamt- X^2 für die Daten der Tabelle 12.8

Variationsursache	X^2	FG	P
Lineare Regression	6,0860	1	0,0142
Abweichungen von der Regression	6,2471	6	0,3999
Insgesamt	12,3331	7	0,0921

Wie die Tabelle zeigt, muß die Hypothese der Unabhängigkeit zugunsten der Hypothese eines linearen Regressionseffektes abgelehnt werden. Die Abweichungen von der Regression können dagegen als zufällig betrachtet werden. Die Anzahl der langen Senkungen steigt somit im Verlauf des Gedichts an, während die Anzahl der kurzen Senkungen abnimmt. Eine mögliche 'Ursache' für diese Tendenz dürfte in der wachsenden inhaltlichen Spannung des Gedichts zu sehen sein (vgl. Kap. 12.3).

Da sich der Wert der Charakteristik 'Anzahl der kurzen bzw. langen Senkungen pro Strophe' im Verlauf des Gedichts statistisch signifikant ändert, ist diese Charakteristik als eine intratextuelle

Stilcharakteristik anzusehen (vgl. Kap.11.2.3). Im Gedicht "Der Totentanz" dagegen ist der Wert dieser Charakteristik weitgehend stabil (vgl. Tabelle 12.6). Die Charakteristik 'Anzahl der kurzen bzw. langen Silben pro Strophe' ist somit auch als subjektive und/oder objektive Stilcharakteristik anzusehen (vgl. Kap. 11.2.2). Um allerdings entscheiden zu können, welcher Typ von Stilcharakteristik tatsächlich vorliegt, müßten noch eine Reihe weiterer Vergleichstexte hinzugezogen werden.

Der Verlauf der Senkungen im "Erkönig" läßt sich wiederum durch Regressionsgeraden beschreiben. Bezeichnen wir die relativen Häufigkeiten der langen Senkungen als p , die der kurzen Senkungen als q und die Variable 'Strophe' als x , erhalten wir die beiden folgenden Geraden:

$$\hat{p} = 0,2031 + 0,0469x \quad (12.30)$$

$$\hat{q} = 0,7969 - 0,0469x \quad (12.31)$$

Die Regressionskoeffizienten dieser Gleichungen besagen, daß der Anteil der langen Senkungen pro Strophe um durchschnittlich 4,69 % zunimmt, während der Anteil der kurzen Senkungen um den gleichen Betrag abnimmt. Die beiden Regressionskoeffizienten sind somit ein Maß für die Stärke der mit Hilfe des COCHRANschen X^2 -Tests festgestellten Tendenz.

Die anhand der beiden Regressionsgeraden ermittelten Schätzwerte sind zusammen mit den beobachteten Werten in der Tabelle 12.10 aufgeführt. Wie aufgrund des relativ hohen Anteils der Abweichungen von der Regression an der Gesamtvariation (vgl. Tabelle 12.9) zu erwarten war, zeigt sich eine nicht unerhebliche Diskrepanz zwischen beobachteten und geschätzten Werten.

Tab. 12.10: Beobachtete und geschätzte Werte für die relativen Häufigkeiten von langen und kurzen Senkungen pro Strophe im "Erk König"

Strophe	lange Senkungen		kurze Senkungen	
	p	$\hat{p}$	q	$\hat{q}$
1	0,2500	0,2500	0,7500	0,7500
2	0,2500	0,2969	0,7500	0,7031
3	0,2500	0,3438	0,7500	0,6562
4	0,4375	0,3906	0,5625	0,6094
5	0,5625	0,4375	0,4375	0,5625
6	0,5000	0,4844	0,5000	0,5156
7	0,6875	0,5313	0,3125	0,4687
8	0,3750	0,5781	0,6250	0,4219

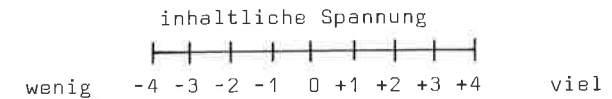
12.3. Korrelation zwischen Inhalt und Formalstruktur

12.3.1. Ich habe bereits angedeutet, daß die im "Erk König" festgestellten formalen Tendenzen möglicherweise mit dem Inhalt des Gedichts korrelieren. Zur Überprüfung dieser Hypothese habe ich, um der im Kap. 3 erhobenen Forderung zu entsprechen, daß bei einer erfahrungswissenschaftlichen Interpretation nicht das eigene subjektive Rezeptionserlebnis, sondern die objektive Untersuchung des Verhaltens beliebiger Rezipienten zugrunde zu legen ist, ein Rezeptionsexperiment durchgeführt, gegen das sich zwar eine Reihe methodologischer Einwände (z.B. hinsichtlich der Inhaltsvalidität oder des Skalenaufbaus) vorbringen lassen, das m.E. jedoch trotz der möglichen Einwände zur Objektivierung der Analyse beiträgt.

Bei diesem Experiment wurde 20 Versuchspersonen (Vpn) der Text des "Erk König" zusammen mit der folgenden Anweisung vorgelegt:⁵

Lesen Sie bitte das Gedicht "Erk König" von GOETHE zuerst einmal insgesamt durch. Dann lesen Sie bitte das Gedicht noch einmal und

stufen die einzelnen Strophen auf der folgenden Skala ein:



12.3.2. Die Daten aus dem Experiment finden sich in Tabelle 12.11. Es soll nun zuerst geprüft werden, ob die einzelnen Strophen als gleich spannend (H_0) oder als unterschiedlich spannend (H_1) beurteilt werden. Da es sich um ordinalskalierte Daten (vgl. Kap. 11.4) aus abhängigen Stichproben handelt, bietet es sich an, eine FRIEDMANN-Rangvarianzanalyse durchzuführen (vgl. FRIEDMANN 1937; BRADLEY 1968:123ff).⁶

Zur Überprüfung des Unterschieds zwischen den Strophen (Spalten) werden die k Werte in jeder Zeile mit Rangzahlen versehen. Treten gleiche Rangplätze auf, d.h. liegen sog. Bindungen (*ties*) vor, werden die entsprechenden Rangplätze gemittelt. Sind alle $k \cdot n$ Stichprobenwerte mit Rangplätzen versehen, werden in jeder Spalte die Rangsummen R_j gebildet. Die in der geschilderten Weise transformierten Daten finden sich in der Tabelle 12.12.

Tab. 12.11: Bewertung der inhaltlichen Spannung der Strophen
des "Erlkönigs" durch 20 Versuchspersonen

Vpn	Strophe							
	1	2	3	4	5	6	7	8
1	-4	-3	-4	-2	-4	-2	0	1
2	0	0	1	1	1	2	3	4
3	-3	2	1	2	2	3	3	4
4	-3	1	0	1	0	2	3	2
5	-2	-1	0	1	2	3	4	4
6	-1	1	1	2	1	2	3	3
7	1	1	1	1	2	2	3	3
8	-4	-3	-3	0	0	1	2	2
9	0	1	0	1	0	1	2	2
10	1	3	1	4	2	4	3	2
11	1	2	0	1	0	1	1	2
12	-4	1	1	1	2	3	4	4
13	1	1	1	1	0	0	1	0
14	-1	1	0	2	1	3	3	2
15	0	2	1	2	1	3	4	2
16	-2	2	0	3	1	4	2	4
17	-4	-1	-4	0	-4	1	1	2
18	0	1	1	1	2	2	3	4
19	-3	-1	0	0	0	1	3	4
20	0	0	0	2	1	2	4	3

Tab. 12.12: Rangzahlen der Zeilen-Werte aus Tab. 12.11

Vpn	Strophe							
	1	2	3	4	5	6	7	8
1	2	4	2	5,5	2	5,5	7	8
2	1,5	1,5	4	4	4	6	7	8
3	1	4	2	4	4	6,5	6,5	8
4	1	4,5	2,5	4,5	2,5	6,5	8	6,5
5	1	2	3	4	5	6	7,5	7,5
6	1	3	3	5,5	3	5,5	7,5	7,5
7	2,5	2,5	2,5	2,5	5,5	5,5	7,5	7,5
8	1	2,5	2,5	4,5	4,5	6	7,5	7,5
9	2	5	2	5	2	5	7,5	7,5
10	1,5	5,5	1,5	7,5	3,5	7,5	5,5	3,5
11	4,5	7,5	1,5	4,5	1,5	4,5	4,5	7,5
12	1	3	3	3	5	6	7,5	7,5
13	6	6	6	6	2	2	6	2
14	1	3,5	2	5,5	3,5	7,5	7,5	5,5
15	1	5	2,5	5	2,5	7	8	5
16	1	4,5	2	6	3	7,5	4,5	7,5
17	2	4	2	5	2	6,5	6,5	8
18	1	3	3	3	5,5	5,5	7	8
19	1	2	4	4	4	6	7	8
20	2	2	2	5,5	4	5,5	8	7
	35	75	53	94,5	69	118	138	137,5

Da in jeder der insgesamt n Zeilen k Ränge vergeben worden sind, kann mit Hilfe der Beziehung

$$\sum_{j=1}^k R_j = n \sum_{x=1}^k x = \frac{nk(k+1)}{2}$$

Überprüft werden, ob die Ränge richtig zugeteilt sind. Wir erhalten

$$\frac{nk(k+1)}{2} = \frac{20(8)9}{2} = 720 = \sum_{j=1}^8 R_j$$

und somit eine Bestätigung der Rangzuteilung.

Der FRIEDMAN-Test - ich betrachte im folgenden nur die für das vorliegende Problem ausreichende χ^2 -Approximation - beruht auf folgender Überlegung:

Würde die Nullhypothese gelten, dürften sich die Rangsummen R_j nur durch Zufall voneinander unterscheiden. Da unter der Nullhypothese für jede Zeile $k!$ verschiedene Rangordnungen mit jeweils gleicher Wahrscheinlichkeit möglich sind, folgen die k Ränge jeder Zeile einer Rechteckverteilung mit dem Erwartungswert

$$E(X) = \frac{\sum_{x=1}^k x}{k} = \frac{k+1}{2}$$

und der Varianz

$$V(X) = E(X^2) - (E(X))^2 = \frac{(k+1)(2k+1)}{6} - \left(\frac{k+1}{2}\right)^2 = \frac{k^2-1}{12}$$

Die Summe R_j der n Ränge in der Spalte j hat folglich eine Verteilung mit dem Erwartungswert $n(k+1)/2$ und der Varianz $n(k^2-1)/12$.

Für $n \rightarrow \infty$ gilt dann nach dem zentralen Grenzwertsatz, daß

$$Z = \frac{R_j - n(k+1)/2}{\sqrt{n(k^2-1)/12}}$$

einer Normalverteilung $N(0,1)$ folgt. Die Summe von k voneinander unabhängigen quadrierten standardisierten Normalvariablen folgt wiederum einer χ^2 -Verteilung mit k Freiheitsgraden. Da jedoch lediglich $k-1$ der k Rangsummen voneinander unabhängig sind, multiplizieren wir $\sum z_i^2$ mit $(k-1)/k$ und betrachten das Resultat als annähernd χ^2 -verteilt mit $k-1$ Freiheitsgraden (vgl. BRADLEY 1968:126f). Die Teststatistik lautet jetzt

$$\chi_R^2 = \frac{k-1}{k} \sum_{j=1}^k \left[\frac{R_j - n(k+1)/2}{\sqrt{n(k^2-1)/12}} \right]^2 = \frac{12}{nk(k+1)} \sum_{j=1}^k \left[R_j - \frac{n(k+1)}{2} \right]^2 \quad (12.32)$$

Da

$$\bar{R} = \frac{1}{k} \sum_{j=1}^k R_j = \frac{1}{k} \frac{nk(k+1)}{2} = \frac{n(k+1)}{2}$$

ist, kann (12.32) auch als

$$\chi_R^2 = \frac{6}{k} \sum_{j=1}^k \frac{(R_j - \bar{R})^2}{\bar{R}} = \frac{6}{\sum_{j=1}^k R_j} \sum_{j=1}^k (R_j - \bar{R})^2 \quad (12.33)$$

geschrieben werden. (12.33) hat zwar den Vorteil, daß der χ^2 -Charakter der Prüfgröße deutlich erkennbar ist, üblicherweise wird jedoch die Formel

$$\chi_R^2 = \left[\frac{12}{nk(k+1)} \sum_{j=1}^k R_j^2 \right] - 3n(k+1) \quad (12.34)$$

verwendet, die man ebenfalls nach einigen Umformungen aus (12.32) erhält.⁷

Für die Daten der Tabelle 12.12 ergibt sich unter Benutzung der Formel (12.33) der folgende χ^2_R -Wert:

$$\chi^2_R = \frac{6}{720} ((75 - 90)^2 + \dots + (137,5 - 90)^2) = 86,87$$

Die Wahrscheinlichkeit, unter der Nullhypothese ein $\chi^2 \geq 86,87$ zu erhalten, ist $\ll 0,0005$. Man kann deshalb annehmen, daß die einzelnen Strophen von den Beurteilern als unterschiedlich spannend eingestuft werden.

12.3.3. Die Stärke der Übereinstimmung der Beurteiler hinsichtlich der Rangordnung der Strophen kann mit Hilfe des KENDALL-schen Konkordanzkoeffizienten W gemessen werden (vgl. KENDALL 1962: 94ff).

Bei der Herleitung dieses Koeffizienten geht man davon aus, daß im Falle einer absoluten Übereinstimmung zwischen den Beurteilern die k Rangsummen aus irgendeiner Anordnung der Werte $n, 2n, \dots, kn$ bestehen würden. Da die mittlere Rangsumme $\bar{R} = n(k+1)/2$ ist, beträgt die Summe der quadrierten Abweichungen der k Rangsummen von der mittleren Rangsumme

$$\sum_{x=1}^k (x_n - \frac{n(k+1)}{2})^2 = \frac{n^2(k^3 - k)}{12}. \quad (12.35)$$

Dies ist der Maximalwert, den die Summe der quadrierten Abweichungen erreichen kann. Den Konkordanzkoeffizienten W erhält man, indem man die Summe der beobachteten quadrierten Abweichungen

$$S = \sum_{j=1}^k (R_j - \frac{n(k+1)}{2})^2 \quad (12.36)$$

durch (12.35) dividiert:

$$W = \frac{\sum_{j=1}^k (R_j - \frac{n(k+1)}{2})^2}{n^2(k^3 - k)} \quad (12.37)$$

Stimmen alle Beurteiler überein, ist S gleich (12.35) und W folglich 1. Stimmen die Beurteiler überhaupt nicht überein, sind die k Rangsummen gleich der mittleren Rangsumme. S und damit auch W ist dann gleich 0.

Für Tabelle 12.12 erhalten wir

$$W = \frac{12(10424,5)}{20^2(8^3 - 8)} = 0,6205.$$

Dieser Wert weist auf eine ziemlich hohe Übereinstimmung zwischen den Beurteilern hin. Dies bedeutet jedoch lediglich, daß die Vpn hinsichtlich der Einstufung der Reihenfolge der einzelnen Strophen auf der Rating-Skala übereinstimmen. In der absoluten Einstufung der Strophen (z.B. eher positive oder eher negative Werte auf der Rating-Skala) können, wie die Tabelle 12.11 zeigt, dennoch erhebliche Unterschiede zwischen den einzelnen Vpn bestehen.

Bei der Berechnung von W muß noch berücksichtigt werden, daß eine größere Anzahl von Bindungen (gleiche Rangplätze) zu einer Verringerung des Werts von W führt. In diesem Fall sollte W folgendermaßen korrigiert werden (vgl. KENDALL 1962:96ff):

$$W_B = \frac{\sum_{j=1}^k (R_j - \frac{n(k+1)}{2})^2}{n^2(k^3 - k) - nT} \quad (12.38)$$

wobei

$$T = \sum_{i=1}^n \sum_{j=1}^{r_i} (t_{ij}^3 - t_{ij})$$

ist und r_i die Anzahl der Bindungen in der i -ten Zeile und t_{ij} die Länge der j -ten Bindung in der i -ten Zeile bezeichnet.

Da der Korrekturfaktor T für die Tabelle 12.12 gleich 736 ist, beträgt der korrigierte Konkordanzkoeffizient

$$W_B = 0,6694.$$

Wie zu erwarten war, ist dieser Wert etwas größer als der entsprechende unkorrigierte Wert.

Zwischen W und X_R^2 besteht die Beziehung

$$X_R^2 = n(k - 1) W. \quad (12.39)$$

Multipliziert man nämlich (12.37) mit $n(k - 1)$ erhält man die Formel (12.32) für X_R^2 .

Die aufgezeigte Beziehung wollen wir dazu benutzen, um eine Formel für die FRIEDMAN-Rangvarianzanalyse im Falle von Bindungen abzuleiten. Wie wir gesehen haben, geht man nämlich bei der Ableitung von X_R^2 davon aus, daß in ein und derselben Rangreihe keine gleichen Rangplätze vorkommen. Berechnet man trotz gleicher Rangplätze die Prüfgröße X_R^2 , kann ein falsches Testresultat die Folge sein, da eine größere Anzahl von Bindungen zu einer Verringerung des Wertes von X_R^2 führt.

Die X_R^2 -Formel für Bindungen ($X_{R,B}^2$) erhalten wir, indem wir in (12.39) für W den korrigierten Konkordanzkoeffizienten W_B (vgl. Formel 12.38) einsetzen. Nach einigen Umformungen ergibt sich dann die folgende Beziehung zwischen X_R^2 und $X_{R,B}^2$:

$$X_{R,B}^2 = \frac{(k^3 - k)n}{(k^3 - k)n - T} X_R^2 \quad (12.40)$$

mit

$$T = \sum_{i=1}^n \sum_{j=1}^{r_i} (t_{ij}^3 - t_{ij})$$

Berechnen wir $X_{R,B}^2$ für die Tabelle 12.12, erhalten wir

$$X_{R,B}^2 = \frac{(8^3 - 8)20}{(8^3 - 8)20 - 736} 86,87 = 93,71.$$

Da wir bereits einen hochsignifikanten X_R^2 -Wert erhalten hatten, ändert der höhere $X_{R,B}^2$ -Wert nichts an der Interpretation des Testresultats. Bei einem nahezu signifikanten Resultat kann die Nichtberücksichtigung von Bindungen jedoch zu einer irrtümlichen Beibehaltung der Nullhypothese führen.

12.3.4. Mit Hilfe der FRIEDMAN-Rangvarianzanalyse können auch Unterschiede zwischen den Versuchspersonen auf Signifikanz geprüft werden. Dazu müssen Rangplätze für die Spalten der Tabelle 12.11 berechnet werden. Die entsprechenden Daten finden sich in Tabelle 12.13. Die Rangvarianzanalyse ergibt den Wert

$$X_R^2 = 71,18.$$

Da bei $FG = 19$ $P(X^2 \geq 46) \ll 0,0005$ ist, kann auf die Berechnung von $X_{R,B}^2$ verzichtet werden. Es besteht somit trotz der hohen Übereinstimmung bei der Beurteilung der Rangordnung der Strophen ein hochsignifikanter Unterschied zwischen den Vpn hinsichtlich der Beurteilung der Spannung der einzelnen Strophen. Der Grund hierfür kann zum einen in individuellen Unterschieden zwischen den Vpn gesehen werden. Zum anderen dürfte jedoch auch die relativ große Vagheit der Kategorie 'inhaltliche Spannung' eine Rolle spielen.

Bei der Überprüfung des Unterschieds zwischen den Strophen haben wir bisher lediglich festgestellt, daß Unterschiede zwischen den Strophen bestehen; wir haben jedoch nicht danach gefragt, welche Strophen sich jeweils unterscheiden.

Diese Frage kann leicht über einen multiplen Vergleich der 8 Rangsummen aus Tabelle 12.12 beantwortet werden (vgl. WILCOXON/WILCOX 1964). Bei diesem Vergleich ermittelt man lediglich die Absolutbeträge der insgesamt $\binom{8}{2}$ Differenzen zwischen den Rangsummen.

Tab. 12.13: Rangzahlen der Spalten-Werte aus Tab. 12.11

Vpn	Strophe								
	1	2	3	4	5	6	7	8	R _i
1	2,5	1,5	1,5	1	1,5	1	1	2	12
2	14	6,5	16	9	11,5	10,5	12	17	96,5
3	6	17,5	16	16	17,5	16	12	17	118
4	6	11,5	7,5	9	5,5	10,5	12	6,5	68,5
5	8,5	4	7,5	9	17,5	16	18,5	17	98
6	10,5	11,5	16	16	11,5	10,5	12	12	100
7	18,5	11,5	16	9	17,5	10,5	12	12	107
8	2,5	1,5	3	3	5,5	5	6	6,5	33
9	14	11,5	7,5	9	5,5	5	6	6,5	65
10	18,5	20	16	20	17,5	19,5	12	6,5	130
11	18,5	17,5	7,5	9	5,5	5	3	6,5	72,5
12	2,5	11,5	16	9	17,5	16	18,5	17	108
13	18,5	11,5	16	9	5,5	2	3	1	66,5
14	10,5	11,5	7,5	16	11,5	16	12	6,5	91,5
15	14	17,5	16	16	11,5	16	18,5	6,5	116
16	8,5	17,5	7,5	19	11,5	19,5	6	17	106,5
17	2,5	4	1,5	3	1,5	5	3	6,5	27
18	14	11,5	16	9	17,5	10,5	12	17	107,5
19	6	4	7,5	3	5,5	5	12	17	60
20	14	6,5	7,5	16	11,5	10,5	18,5	12	96,5
	1680								

Die kritischen Werte für die Differenzen findet man z.B. bei WILCOXON/WILCOX (1964:36ff) oder bei MCDONALD/THOMPSON (1967:490ff) oder auch bei SACHS (1974:427ff). Für $k=8$ und $n=20$ betragen die kritischen Werte 43,1 ($P=0,10$), 47 ($P=0,05$) und 54,6 ($P=0,01$). Die Differenzen zwischen den Rangsummen sind in der Tabelle 12.14 aufgeführt. Signifikante Differenzen sind dabei durch * gekennzeichnet.

Wie die Tabelle zeigt, unterscheiden sich vor allem die Strophen 4, 6, 7 und 8 von den übrigen Strophen. Besonders deutlich ist der Unterschied bei den Strophen 7 und 8 ausgeprägt: Der Unterschied zu den Strophen 1, 2, 3 und 5 ist hochsignifikant ($P = 0,01$), der Unterschied zu Strophe 4 ist signifikant mit $P = 0,10$. Lediglich zur Strophe 6 besteht kein Unterschied. Der Vergleich der Rangsummen zeigt somit deutlich, daß für die Vpn die inhaltliche Spannung in der zweiten Gedichthälfte ihren Höhepunkt erreicht. Allerdings ist bei diesem Ergebnis zu berücksichtigen, daß möglicherweise die Variable 'Reihenfolge der Strophen' das Urteil der Vpn mitbeeinflusst hat und daß deshalb der durchgeführte Test nur beschränkt valide ist (vgl. Kap. 11.3).

Tab. 12.14: Multiple Vergleiche zwischen den Rangsummen aus Tab. 12.12

Strophen	R _j	Strophe						
		2	3	4	5	6	7	8
		75	53	94,5	69	118	138	137,5
1	35	40	18	59,5***	34	83***	113***	112,5***
2	75		22	19,5	6	43*	63***	62,5***
3	53			41,5	16	65***	85***	84,5***
4	94,5				25,5	23,5	43,5*	43*
5	69					49**	69***	68,5***
6	118						20	19,5
7	138							0,5

* signifikant mit $P = 0,10$

** signifikant mit $P = 0,05$

*** signifikant mit $P = 0,01$

12.3.5. Abschließend sollen nun die Daten zur inhaltlichen Spannung aus Tabelle 12.12 mit der formalen Charakteristik 'Anzahl der langen Senkungen' (vgl. Tab. 12.8) korreliert werden.

Ordnet man den Häufigkeiten der langen Senkungen aus Tabelle 12.8 Rangzahlen zu, erhält man die Rangreihe

2 2 2 5,5 6 5,5 7 4.

Zur Berechnung der Korrelation dieser Rangreihe mit den Rangreihen aus Tabelle 12.12 benutzen wir den SPEARMANSchen Rangkorrelationskoeffizienten r_S (zur Herleitung von r_S vgl. z.B. WOLF 1974:220ff).

Der Korrelationskoeffizient r_S ist definiert als

$$r_S = 1 - \frac{6 \sum_{i=1}^n d_i^2}{n^3 - n}, \quad (-1 \leq r_S \leq 1) \quad (12.41)$$

wobei d_i die einzelnen Rangplatzdifferenzen bezeichnet. Im Fall einer größeren Anzahl von Bindungen in ein und derselben Rangreihe sollte r_S durch einen Korrekturfaktor für Bindungen korrigiert werden. Der korrigierte Rangkorrelationskoeffizient $r_{S,B}$ lautet

$$r_{S,B} = \frac{(n^3 - n) - \frac{1}{2}(T_x + T_y) - 6 \sum_{i=1}^n d_i^2}{\sqrt{(n^3 - n - T_x)(n^3 - n - T_y)}} \quad (12.42)$$

mit

$$T_x = \sum (t_x^3 - t_x) \text{ und } T_y = \sum (t_y^3 - t_y),$$

wobei T_x gleich der Anzahl der Bindungen gleicher Ranggröße in der Rangreihe x und T_y gleich der Anzahl der Bindungen gleicher Ranggröße in der Rangreihe y ist.

Die Korrektur für Bindungen hat zur Konsequenz, daß der Wert eines positiven Korrelationskoeffizienten r_S verringert, der Wert eines negativen Korrelationskoeffizienten dagegen vergrößert wird.

Wir berechnen nun $r_{S,B}$ für die Vp 1 der Tabelle 12.12. Die zu vergleichenden Rangreihen und die Summe der quadrierten Rangplatzdifferenzen lauten

x:	2	2	2	5,5	6	5,5	7	4
y:	2	4	2	5,5	2	5,5	7	8
d_i^2 :	0	4	0	0	16	0	0	16

Da $T_x = 3^3 - 3 + 2^3 - 2 = 30 = T_y$ ist, erhalten wir

$$r_{S,B} = \frac{(8^3 - 8) - \frac{1}{2}(30 + 30) - 6(36)}{\sqrt{(8^3 - 8 - 30)(8^3 - 8 - 30)}} = 0,5443.$$

Dieser Wert ist etwas niedriger als der entsprechende unkorrigierte Wert, der 0,5714 beträgt.

Die exakten Überschreitungswahrscheinlichkeiten für $r_{S,B}$ lassen sich mit Hilfe der Tafeln der Unterschreitungswahrscheinlichkeiten für Σd_i^2 von OWEN (1962:400ff) berechnen. Für $r_{S,B} = 0,5443$ beträgt die korrigierte Summe der Rangplatzdifferenzen 38,28 und die entsprechende Wahrscheinlichkeit 0,087.

Die Korrelationskoeffizienten sind zusammen mit den Überschreitungswahrscheinlichkeiten in der Tabelle 12.15 aufgeführt.

Von den 20 Korrelationskoeffizienten sind bei einem Signifikanzniveau von $P = 0,05$ 13 signifikant. Lediglich bei 2 Vpn ist die Korrelation negativ, wobei allerdings der Wert noch im Annahmebereich der Nullhypothese liegt. Auf Befragen gaben diese beiden Vpn an, daß ihre Einstellung gegenüber Lyrik negativ sei. Die Einstellung der übrigen Vpn zur Lyrik war dagegen fast ausschließlich mehr oder minder positiv.

Tab. 12.15: Korrelation zwischen der inhaltlichen Spannung und der Anzahl der langen Senkungen im "Erlkönig"

Vpn	$r_{S,B}$	P	Vpn	$r_{S,B}$	P
1	0,5443	0,087	11	-0,0265	0,536
2	0,6582	0,044	12	0,7089	0,031
3	0,6392	0,050	13	-0,1743	0,675
4	0,6375	0,050	14	0,8001	0,012
5	0,7780	0,016	15	0,6582	0,044
6	0,6879	0,036	16	0,5156	0,101
7	0,7027	0,032	17	0,5253	0,097
8	0,7813	0,015	18	0,6772	0,039
9	0,5004	0,108	19	0,6541	0,045
10	0,6195	0,057	20	0,8228	0,009

Wir wollen nun die 20 Testresultate zu einem Gesamttest verbinden (vgl. KAPUR/SAXENA 1970:377ff). Nimmt man an, daß für jede einzelne Überschreitungswahrscheinlichkeit P_i

$$P_i = P(\chi^2 \geq \chi_i^2) = \frac{1}{2} \int_{\chi_i^2}^{\infty} e^{-\frac{1}{2} \chi^2} d\chi^2 = e^{-\frac{1}{2} \chi_i^2}$$

gilt, dann folgt, daß

$$\chi_i^2 = -2 \ln P_i = 2 \ln \frac{1}{P_i}$$

ist.

Für die dritte Vp z.B. ist $P_3 = 0,05$. Die Transformation ergibt

$$-2 \ln 0,05 = 5,99 = \chi_3^2$$

Wie die χ^2 -Tafeln zeigen, entspricht diesem Wert eine Wahrscheinlichkeit von 0,05.

Da für n Überschreitungswahrscheinlichkeiten aufgrund der Additivität der χ^2 -Variablen

$$\chi^2 = \chi_1^2 + \chi_2^2 + \dots + \chi_n^2 = -2 \ln(P_1 P_2 \dots P_n)$$

mit $FG = 2n$ gilt, erhalten wir für die 20 Wahrscheinlichkeiten der Tabelle 12.15

$$\chi^2 = -2 \ln(0,087 \cdot 0,044 \dots 0,009) = 118,93$$

mit $FG = 40$.

Bei 40 Freiheitsgraden ist $P(\chi^2 \geq 76,1) = 0,0005$. Es kann somit angenommen werden, daß insgesamt gesehen eine hochsignifikante Korrelation zwischen inhaltlicher Spannung und langen Senkungen besteht. Die Frage, ob eine solche Korrelation zwischen Formalstruktur und Inhalt vom Autor intendiert ist oder ob eine unterbewußte sprachliche Gestaltung vorliegt (vgl. zu diesem Problem JAKOBSON 1971), soll im Rahmen dieser Arbeit nicht weiter verfolgt werden. Hierzu wären u.a. biographische und stilpsychologische Untersuchungen notwendig (vgl. Kap. 3).⁸

12.4. Verstyp und Verslänge

12.4.1. Wir hatten festgestellt, daß im "Erlkönig" die Anzahl der Senkungen zwischen zwei Hebungen z.B. im Verhältnis zu "Der Totentanz" relativ variabel ist. Dies hat zur Folge, daß der "Erlkönig" auch eine verhältnismäßig hohe Anzahl verschiedener Verstypen aufweist.

Die verschiedenen Verstypen finden sich geordnet nach der Anzahl der unmarkierten Silben in der Tabelle 12.16.

Tab. 12.16: Häufigkeiten der verschiedenen Verstypen im "Erlkönig"

Typ	U	Notation	Zeile	f_i	%
1	4	UMUMUMUM	8,9	2	6,250
2	5	UMUUMUMUM	1	1	3,125
3	5	UMUMUUMUM	2,4,7,17,32	5	15,625
4	5	UMUMUMUUM	3,6,10,11,15,16,23,31	8	25,000
5	5	MUUMUUMUM	22,28	2	6,250
6	6	UUMUMUUMUM	18	1	3,125
7	6	UUMUMUMUUM	12	1	3,125
8	6	UMUUMUMUUM	24	1	3,125
9	6	UMUMUUMUUM	5,14,29,30	4	12,500
10	7	UUMUMUUMUUM	19	1	3,125
11	7	UMUUMUUMUUM	13,20,21,26,27	5	15,625
12	8	UMUUMUUMUUM	25	1	3,125

Wie aufgrund der Analyse der Senkungen zu erwarten war, zeigt sich eine deutliche Vorliebe für Verse mit langen Senkungen am Versende. Sieht man die Typen 5 (fehlende Senkung am Versanfang) und 12 (trissyllabische Senkung) als 'Ausnahmen' an, dann findet sich vor jeder Hebung entweder eine mono- oder eine bisyllabische Senkung. Die theoretisch mögliche Anzahl der Verstypen beträgt unter dieser Voraussetzung $2^4 = 16$. Von diesen 16 Möglichkeiten sind jedoch nur 10 realisiert. Vier der insgesamt 12 Typen stellen allein 68,75 % aller Verse. Auf den Typ 4 entfallen sogar 25 %. Der Autor scheint also einige wenige Typen zu bevorzugen.

Um zu überprüfen, ob diese Tendenz dem Zufall zuzuschreiben ist, testen wir die Häufigkeiten der Verstypen auf Gleichverteilung. Bei Annahme einer Gleichverteilung beträgt der Erwartungswert (E) der Häufigkeiten $32/12 = 2,67$. Da der zur Prüfung auf Gleichverteilung üblicherweise verwendete χ^2 -Test bei einer größeren Zahl von Erwartungswerten kleiner als 5 nicht benutzt werden sollte, verwenden wir zur Überprüfung der Hypothese die auf KULLBACK u.a. zurückgehende Informationsstatistik (*minimum discrimination information statistic*).

Bei kategorischen Daten lautet die informationsstatistische Prüfgröße (vgl. KULLBACK u.a. 1962:218f)

$$2I = 2 \sum_{i=1}^k O_i \ln \frac{O_i}{E_i} = 2 \sum_{i=1}^k f_i \ln \frac{kf_i}{n} \quad (12.43)$$

Entwickelt man in (12.43) $\ln(O_i/E_i)$ in eine Taylorreihe und vernachlässigt die Glieder der Ordnung $O(n^{-1/2})$, erhält man aus (12.43) das übliche χ^2 -Testkriterium. $2I$ ist somit annähernd χ^2 -verteilt mit $k - 1$ Freiheitsgraden.

Für den Test auf Gleichverteilung kann (12.43) auch als

$$2I = 2 \sum_{i=1}^k f_i \ln \frac{kf_i}{n} \quad (12.44)$$

oder als

$$2I = \sum_{i=1}^k 2f_i \ln f_i - 2n \ln n + 2n \ln k \quad (12.45)$$

geschrieben werden. (12.45) ist rechnerisch etwas einfacher und hat den Vorteil, daß für $2f_i \ln f_i$ umfangreiche Tafeln zur Verfügung stehen (z.B. KULLBACK u.a. 1962 für $f_i = 1$ bis 10000 und BLÖSCHL 1966 für $f_i = 1$ bis 5000).

Mit $k = 12$ und $n = 32$ erhalten wir für die Häufigkeiten aus Tabelle 12.16

$$2I = 19,32.$$

Für $FG = 11$ ist $\chi^2_{0,05} = 19,7$. Die Abweichung von der Gleichverteilung ist somit nahezu signifikant. Legt man in Anbetracht des geringen Stichprobenumfangs das weniger konservative Signifikanzniveau von $P = 0,10$ zugrunde, dann könnte die Hypothese, daß keine Bevorzugung bestimmter Verstypen vorliegt, abgelehnt werden.

Eine Tendenz zur Verwendung einiger weniger Verstypen hat auch FÜCKS (1968:77) in GOETHEs "Torquato Tasso" und in SCHILLERs "Wallenstein" festgestellt. Im "Torquato Tasso" fanden sich in einer Stichprobe von 500 Versen 16 Verstypen. Sechs der Typen machten aber 97,4 % der Verse aus. Eine etwas breitere Verteilung zeigte der "Wallenstein".

12.4.2. Als Maß für die metrische Variation eines Textes kann die Wiederholungsrate (*repeat rate*) benutzt werden, die als

$$R = \sum_{i=1}^k p_i^2 \quad (12.46)$$

definiert wird (vgl. HERDAN 1962:36ff, 1966:271ff). R ist ein Maß der Gleichverteilung,⁹ das umso größer wird, je ähnlichere Häufigkeiten die Verstypen aufweisen. R wird 1, wenn ein Verstyp die Wahr-

scheinlichkeit 1 hat, d.h. wenn ein Text nur einen einzigen Verstyp aufweist. Wenn alle $\hat{p}_i$ gleich $1/k$ sind, erreicht R sein Minimum $1/k$.

Für den "Erlkönig" ist $R = 0,1406$. In "Der Totentanz", wo sich lediglich 3 Verstypen mit den relativen Häufigkeiten $2/49$, $21/49$ und $26/49$ finden, beträgt R dagegen $0,4669$.

12.4.3. Für den Rhythmus eines metrischen Textes ist neben der Anzahl der verschiedenen Verstypen auch die Anzahl und Verteilung der Wortgrenzen im Vers von Bedeutung (vgl. Kap. 5.4).

Wir wollen nun zuerst prüfen, ob zwischen "Erlkönig" und "Der Totentanz" hinsichtlich der Anzahl der Wörter pro Vers und damit auch hinsichtlich der Verteilung der Wortgrenzen im Vers ein Unterschied besteht. Die Tabelle 12.17 zeigt die Verteilung der Zahl der Wörter pro Vers im "Erlkönig" und in "Der Totentanz".

Tab. 12.17: Anzahl der Wörter pro Vers im "Erlkönig" und in "Der Totentanz"

Wörter	Verse			
	Erlkönig		Totentanz	
	f_i	%	f_i	%
3	-	-	1	2,04
4	-	-	2	4,08
5	2	6,25	5	10,20
6	10	31,25	9	18,37
7	9	28,13	14	28,57
8	7	21,88	11	22,45
9	4	12,50	6	12,24
10	-	-	-	-
11	-	-	1	2,04

Wir berechnen zuerst einige für die weitere Analyse wichtige Kennwerte. Die Schätzfunktionen der Anfangsmomente, der Zentralmomente, des Variationskoeffizienten, der Spannweite, der Schiefe und der Wölbung sind im Fall diskreter Zufallsvariablen folgendermaßen definiert:

(1) Anfangsmomente

$$m'_r = \frac{1}{n} \sum_i f_i x_i^r \quad (12.47)$$

$$m'_1 = \bar{x} = \frac{1}{n} \sum_i f_i x_i \quad (\text{Mittelwert}) \quad (12.48)$$

$$m'_2 = \frac{1}{n} \sum_i f_i x_i^2$$

usw.

(2) Zentralmomente

$$m_r = \frac{1}{n} \sum_i f_i (x_i - \bar{x})^r \quad (12.49)$$

Varianz:

$$m_2 = S^2 = \frac{1}{n} \sum_i f_i (x_i - \bar{x})^2 \quad (12.50)$$

bzw. als erwartungstreue Schätzung

$$s^2 = \frac{1}{n-1} \sum_i f_i (x_i - \bar{x})^2 \quad (12.51)$$

Standardabweichung:

$$S = \sqrt{S^2} \text{ bzw. } s = \sqrt{s^2} \quad (12.52)$$

höhere Zentralmomente:

$$m_3 = \frac{1}{n} \sum_i f_i (x_i - \bar{x})^3 = m'_3 - 3m'_1 m'_2 + 2m'_1{}^3 \quad (12.53)$$

$$m_4 = \frac{1}{n} \sum_i f_i (x_i - \bar{x})^4 = m'_4 - 4m'_1 m'_3 + 6m'_1{}^2 m'_2 - 3m'_1{}^4 \quad (12.54)$$

(3) Variationskoeffizient

$$V = S/\bar{x} \quad (12.55)$$

(4) Spannweite

$$R = x_{\max} - x_{\min} \quad (12.56)$$

(5) Schiefe

$$g_1 = m_3/S^3 \quad (12.57)$$

m_3 ist bei symmetrischen Verteilungen gleich 0. Es gilt deshalb:

$g_1 = 0$	Symmetrie
$g_1 > 0$	positive Schiefe, Rechtsasymmetrie
$g_1 < 0$	negative Schiefe, Linksasymmetrie

(6) Wölbung (Exzeß, Kurtosis)

$$g_2 = \frac{m_4}{S^4} - 3 \quad (12.58)$$

Da die Wölbung der Normalverteilung 3 beträgt, gilt für g_2 :

$g_2 = 0$	normale Wölbung (mesokurtische Verteilung)
$g_2 > 0$	positive Wölbung, steilgipflige (leptokurtische) Verteilung
$g_2 < 0$	negative Wölbung, flachgipflige (platykurtische) Verteilung

Für die Daten der Tabelle 12.17 erhalten wir folgende Werte:

Tab. 12.18: Kenngrößen der Daten aus Tabelle 12.17

	Erlkönig	Totentanz		Erlkönig	Totentanz
$\bar{x}$	7,0313	6,9592	s^2	1,2803	2,3249
m_2'	50,7188	50,7551	s^2	1,3216	2,3733
m_3'	374,9063	384,9592	m_3	0,2862	-0,6132
m_4'	2835,4688	3021,6122	m_4	3,4880	17,6318
g_1	0,1976	-0,1730	V	0,1609	0,2191
g_2	-0,8721	0,2621	R	4,0000	8,0000

Da der beobachtete Unterschied zwischen den Mittelwerten von "Erlkönig" und von "Der Totentanz" sehr klein und deswegen von geringer inhaltlicher Aussagekraft ist, verzichten wir auf einen Signifikanztest und betrachten die beiden Mittelwerte als gleich. Interessanter ist jedoch der Unterschied zwischen den Varianzen. Wie zu erwarten war, ist die Varianz in "Der Totentanz", wo neben 4-hebigen Versen auch 3-hebige Verse vorkommen, erheblich größer als im "Erlkönig".

Wir wollen nun mit Hilfe eines F-Tests prüfen, ob sich die Varianzen beider Gedichte signifikant unterscheiden. Bei der Anwendung des F-Tests ergibt sich allerdings die Schwierigkeit, daß dieser Test normalverteilte Stichproben voraussetzt. Wie jedoch z.B. die Werte für g_1 und g_2 aus Tabelle 12.18 zeigen, scheinen in beiden Gedichten gewisse Abweichungen von der Normalverteilung vorzuliegen. So ist im "Erlkönig" die Verteilung der Anzahl der Silben pro Vers eher rechtsasymmetrisch und flachgipflig, in "Der Totentanz" dagegen eher linksasymmetrisch und steilgipflig. Wie man leicht nachweisen kann, spielt die Normalitätsannahme jedoch gerade bei

Tests zur Überprüfung von Varianzen eine wichtige Rolle. Bei Tests zur Überprüfung von Mittelwerten ist die Normalitätsannahme dagegen weit weniger bedeutsam. Dies läßt sich an den entsprechenden Standardfehlern zeigen. Der Standardfehler der Varianz beträgt

$$\sigma_s^2 = \sqrt{\frac{\mu_4 - \mu_2^2}{n} + \frac{2 \mu_2^2}{n(n-1)}}.$$

Dieser Ausdruck kann nach einigen Umformungen auch als

$$\sigma_s^2 = \sigma^2 \sqrt{\frac{2}{n-1} + \frac{\gamma_2}{n}}$$

geschrieben werden, wobei

$$\gamma_2 = \frac{E(X - E(X))^4}{\sigma^4} - 3$$

der Wölbungskoeffizient der Grundgesamtheit ist.

Bei normalverteilten Grundgesamtheiten ist $\gamma_2 = 0$. Ist nun bei einem Varianztest γ_2 in Wirklichkeit größer als 0, kann dies zu einer irrtümlichen Ablehnung der Nullhypothese führen (Fehler erster Art). Ist $\gamma_2 < 0$ kann ein konservatives Testresultat (Fehler zweiter Art) die Folge sein. Der Standardfehler des Mittelwertes ist dagegen

$$\sigma_{\bar{x}} = \sigma / \sqrt{n},$$

gleichgültig, welche Verteilungsform die Grundgesamtheit hat. Mittelwerttests sind deshalb im Gegensatz zu Varianztests relativ unempfindlich gegenüber Abweichungen von der Normalverteilung.

Es dürfte somit ratsam sein, die Normalität der beiden Verteilungen aus Tabelle 12.17 zu überprüfen. Wir beschränken uns dabei auf die Überprüfung von Schiefe und Wölbung. Die Verteilung der Schiefe und der Wölbung ist relativ ausführlich bei GEBHARDT (1966) tabelliert. Kritische Schranken (5 % und 1 %) geben auch

PEARSON/HARTLEY (1970:207f). Bei $n = 32$, d.h. im Fall des "Erlkönig", beträgt die Überschreitungswahrscheinlichkeit für die Schiefe mehr als 25 % und die Unterschreitungswahrscheinlichkeit für die Wölbung etwa 10 %. Bei $n = 49$, d.h. im Fall von "Der Totentanz", beträgt die Unterschreitungswahrscheinlichkeit für die Schiefe ebenfalls mehr als 25 % und die Überschreitungswahrscheinlichkeit für die Wölbung etwa 20 %. Man kann deshalb hinsichtlich Schiefe und Wölbung die Anzahl der Wörter pro Vers in beiden Gedichten als normalverteilt ansehen. Die Anwendung des F-Tests scheint somit durchaus gerechtfertigt.

Der F-Test beruht auf folgender Überlegung: Bei gleichen Varianzen würde man erwarten, daß der Quotient $F = s_1^2/s_2^2$ einen Wert von ungefähr 1 annimmt, bei ungleichen Varianzen dagegen einen von 1 abweichenden Wert. Da nun Abweichungen von 1 allein durch Zufall bedingt sein können, benötigen wir, um den Unterschied zweier Varianzen auf Signifikanz testen zu können, die Verteilung des F-Quotienten. Diese erhält man folgendermaßen: Sind X und Y unabhängige normalverteilte Zufallsvariablen mit den Varianzen σ_1^2 und σ_2^2 , dann gilt

$$\chi_1^2 = \frac{(n_1 - 1)s_1^2}{\sigma_1^2} \quad \text{und} \quad \chi_2^2 = \frac{(n_2 - 1)s_2^2}{\sigma_2^2}$$

mit $n_1 - 1$ und $n_2 - 1$ Freiheitsgraden. Der Quotient zweier χ^2 -Variablen, dividiert durch die Anzahl der Freiheitsgrade, folgt nun, wie R.A. FISHER gezeigt hat, einer F-Verteilung mit $(n_1 - 1, n_2 - 1)$ Freiheitsgraden. Es gilt somit

$$\frac{s_1^2/s_2^2}{\sigma_1^2/\sigma_2^2} = \frac{\chi_1^2/(n_1-1)}{\chi_2^2/(n_2-1)} = \frac{\sigma_1^2}{\sigma_2^2} F_{n_1-1, n_2-1}$$

Nimmt man unter der Nullhypothese an, daß sich die Varianzen der zu vergleichenden Grundgesamtheiten nicht unterscheiden, folgt $\sigma_1^2 = \sigma_2^2$, und man erhält die gesuchte Teststatistik:

$$F_{n_1-1, n_2-1} = \frac{s_1^2}{s_2^2} \quad (12.59)$$

Wir testen nun, ob die Varianz der Anzahl der Silben pro Vers in "Der Totentanz" größer als im "Erlkönig" ist. Der Vergleich zwischen den beiden Varianzen ergibt $F_{48,31} = 2,3733/1,3216 = 1,7958$. Interpoliert man linear in der F-Tabelle, erhält man für $P = 0,05$ einen kritischen Wert von 1,786. Die Hypothese $\sigma_1^2 = \sigma_2^2$ kann somit zugunsten der Alternativhypothese $\sigma_1^2 > \sigma_2^2$ verworfen werden. Die Ursache für die größere Variation in "Der Totentanz" dürfte dabei, wie bereits erwähnt wurde, in der andersartigen metrischen Gestaltung zu suchen sein.

Wir wollen nun prüfen, ob zwischen den einzelnen Strophen des "Erlkönig" ein Unterscheid in der mittleren Anzahl der Wörter pro Vers besteht. Diese Fragestellung kann mit Hilfe einer parametrischen Varianzanalyse untersucht werden. Bei diesem Verfahren wird jedoch vorausgesetzt, daß die Varianzen der zu vergleichenden Stichproben homogen sind (Homoskedastizitätsbedingung) und daß die einzelnen Stichproben zumindest annähernd normalverteilt sind. Wie wir gezeigt haben, sind Mittelwertvergleiche jedoch verhältnismäßig unempfindlich gegenüber Abweichungen von der Normalverteilung. In vielen Fällen kann deshalb auf die Überprüfung der Normalverteilungsannahme verzichtet werden. Was die Homoskedastizitätsbedingung betrifft, so wird bei ungleichen Stichprobengrößen die Homogenität von mehr als zwei Varianzen meist mit dem BARTLETT-Test geprüft. Dieser Test ist jedoch äußerst empfindlich gegenüber Abweichungen von der Normalität. Es sollte deshalb in solchen Fällen ein robusteres Verfahren benutzt werden (vgl. den Übersichtsartikel von GAMES/WINKLER/PROBERT 1972). Da bei gleichen Stichprobengrößen der varianzanalytische Mittelwertvergleich jedoch auch gegenüber Abweichungen von der Varianzhomogenität relativ unempfindlich ist, können wir im vorliegenden Fall auf eine vorhergehende Überprüfung der Varianzgleichheit verzichten.

Beim varianzanalytischen Vergleich der mittleren Anzahl der Wörter pro Vers in den Strophen des "Erlkönig" legen wir das folgende Modell zugrunde. Wir betrachten die 8 Strophen als 8 unabhängige Zufallsstichproben aus normalverteilten Grundgesamtheiten. Geprüft wird die Nullhypothese, daß die 8 Strophenmittelwerte aus ein und derselben Grundgesamtheit mit dem Mittelwert μ' stammen, d.h. daß $H_0: \mu'_1 = \mu'_2 = \dots = \mu'_8 = \mu'$ gilt.

Faßt man die k Stichproben zu einer Gesamtstichprobe zusammen, so hat diese die Varianz

$$s_{ges}^2 = \frac{1}{n-1} \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x})^2 \quad (12.60)$$

mit

$$n = \sum_{i=1}^k n_i \text{ und } \bar{x} = \frac{1}{n} \sum_{i=1}^k \sum_{j=1}^{n_i} x_{ij}.$$

Die Varianzanalyse beruht auf einer Zerlegung der Summe der Abweichungsquadrate der Gesamtstichprobe

$$SAQ_{ges} = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x})^2 \quad (12.61)$$

in die Summe der Abweichungsquadrate innerhalb der Stichproben

$$SAQ_{in} = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_i)^2 \quad (12.62)$$

und in die Summe der Abweichungsquadrate zwischen den Stichproben

$$SAQ_{zwi} = \sum_{i=1}^k n_i (\bar{x}_i - \bar{x})^2. \quad (12.63)$$

Es gilt somit (vgl. z.B. BRUNK 1965:286ff)

$$SAQ_{ges} = SAQ_{in} + SAQ_{zwi}$$

mit $FG = n - 1$ für SAQ_{ges} , $FG = n - k$ für SAQ_{in} und $FG = k - 1$ für SAQ_{zwi} .

Dividiert man die Summen der Abweichungsquadrate durch die entsprechenden Freiheitsgrade, erhält man die sog. mittleren Abweichungsquadrate MAQ_{ges} , MAQ_{in} und MAQ_{zwi} , die auch als Gesamtvarianz, Binnenvarianz und Zwischenvarianz bezeichnet werden. Unter der Nullhypothese sind MAQ_{ges} , MAQ_{in} und MAQ_{zwi} jeweils erwartungstreue Schätzfunktionen der unbekannten Populationsvarianz σ^2 . Bildet man nun analog zum Vorgehen bei der Bildung der Prüfstatistik für den Vergleich zweier Varianzen den Quotienten MAQ_{zwi}/MAQ_{in} , so würde man erwarten, daß dieser Quotient unter der Nullhypothese gleich 1 ist. Liegt jedoch ein Bedingungseffekt vor, d.h. übt im vorliegenden Fall die unabhängige Variable 'Strophe' einen Einfluß aus, würde man erwarten, daß MAQ_{zwi} größer als MAQ_{in} ist und daß damit auch der Quotient größer als 1 ist. Es gilt nun unter der Nullhypothese (vgl. z.B. BRUNK 1965:231f)

$$\frac{MAQ_{zwi}}{\sigma^2} = \frac{\chi^2}{FG_{zwi}} \text{ und } \frac{MAQ_{in}}{\sigma^2} = \frac{\chi^2}{FG_{in}}$$

mit $FG_{zwi} = k - 1$ und $FG_{in} = n - k$.

Der Quotient dieser beiden Ausdrücke ist

$$\frac{MAQ_{zwi}/\sigma^2}{MAQ_{in}/\sigma^2} = \frac{MAQ_{zwi}}{MAQ_{in}} = \frac{s_{zwi}^2}{s_{in}^2} = \frac{\chi^2/FG_{zwi}}{\chi^2/FG_{in}}.$$

Da der Quotient zweier unabhängiger χ^2 -Variablen, dividiert durch die Zahl der Freiheitsgrade, wiederum einer F-Verteilung entspricht, erhält man schließlich die Prüfstatistik

$$\hat{F} = \frac{s_{\text{zwi}}^2}{s_{\text{in}}^2} = \frac{\frac{1}{k-1} \sum_{i=1}^k n_i (\bar{x}_i - \bar{x})^2}{\frac{1}{n-k} \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_i)^2} \quad (12.64)$$

Die Tabelle 12.19 zeigt die Anzahl der Wörter pro Vers in den einzelnen Strophen des "Erlkönig".

Tab. 12.19: Anzahl der Wörter pro Vers in den einzelnen Strophen des "Erlkönig"

Vers	Strophe							
	1	2	3	4	5	6	7	8
1	8	9	7	8	7	9	8	6
2	7	6	7	5	6	5	9	7
3	8	6	7	6	6	8	9	7
4	8	6	6	6	8	7	6	7
$\bar{x}_i$	7,75	6,75	6,75	6,25	6,75	7,25	8	6,75

Die verschiedenen Varianzkomponenten sind in der Tabelle 12.20 aufgeführt.

Tab. 12.20: Varianzanalyse der Daten aus Tabelle 12.19

Variabilität	SAQ	FG	MAQ
zwischen	9,72	7	1,3884
innerhalb	31,25	24	1,3021
gesamt	40,97	31	1,3216

Da MAQ_{zwi} und MAQ_{in} fast gleich sind, beträgt $\hat{F}$ lediglich $1,3884/1,3021 = 1,0663$. Dieser Wert bleibt weit unter dem kritischen Wert $F_{7,24;0,05} = 2,43$. Es besteht somit kein Unterschied zwischen den Strophen des "Erlkönig" hinsichtlich der mittleren Anzahl der Wörter bzw. Wortgrenzen pro Vers. Untersucht man anstelle der Variablen 'Strophe' den Einfluß der Variablen 'Gedichtpartie', ergibt sich ebenfalls kein Unterschied zwischen den Mittelwerten der Erzähl-, Vater-, Sohn- und Erlkönigpartie. Wählt man anstelle der abhängigen Variablen 'Anzahl der Wörter pro Vers' die Variable 'Anzahl der Silben pro Vers', erhält man gleichfalls ein nicht-signifikantes Ergebnis. Die untersuchten Variablen können somit - im Gegensatz z.B. zur Variablen 'Anzahl der Senkungen in den Strophen 1 bis 4 und 5 bis 8' - nicht als intratextuelle Stilcharakteristiken des "Erlkönig" angesehen werden.

12.4.4. Bisher haben wir die nach Kriterien wie Senkungstyp oder Anzahl der Wörter klassifizierten Verse des "Erlkönig" lediglich anhand der Häufigkeit ihres Auftretens charakterisiert. Wie jedoch in Kap. 4 gezeigt worden ist, muß zur Charakterisierung einer Zufalls-Folge auch die Anordnung der Elemente der Folge berücksichtigt werden.

Die Tabelle 12.21 zeigt die Anzahl der Silben pro Vers im "Erlkönig".

Tab. 12.21: Anzahl der Silben (S) pro Vers (V) im "Erlkönig"

V	S	V	S	V	S	V	S
1	9	9	8	17	9	25	12
2	9	10	9	18	10	26	11
3	9	11	9	19	11	27	11
4	9	12	10	20	11	28	9
5	10	13	11	21	11	29	10
6	9	14	10	22	9	30	10
7	9	15	9	23	9	31	9
8	8	16	9	24	10	32	9

Es fällt sofort auf, daß lediglich Verse mit 8, 9, 10, 11 oder 12 Silben vorkommen. Außerdem scheinen auch signifikante Unterschiede hinsichtlich der Häufigkeiten der einzelnen Verstypen vorzuliegen.

Die Häufigkeiten sind in der Tabelle 12.22 zusammengestellt.

Tab. 12.22: Zusammenfassung der Daten aus Tabelle 12.21

Silben	8	9	10	11	12
Häufigkeiten	2	16	7	6	1

Führt man wiederum einen Test auf Gleichverteilung durch, erhält man einen χ^2 -Wert von 22,06, der den kritischen Wert von $\chi^2_{4,0,05} = 9,49$ weit überschreitet. Die Hypothese einer Zufallsfolge mit $p_1 = p_2 = \dots = p_k$ kann somit nicht aufrecht erhalten werden.

Wie die Tabelle 12.21 zeigt, scheint jedoch auch bei der Anordnung der Verstypen (V) eine Tendenz zur Strukturierung vorzuliegen. Es soll jetzt untersucht werden, ob der "Erlkönig" - aufgefaßt als Folge der Elemente (Grundsymbole) V_8, V_9, V_{10}, V_{11} und V_{12} (vgl. Kap. 4) - auch hinsichtlich der Anordnung dieser Elemente von einer Zufallsfolge abweicht.¹⁰

Zur Beantwortung dieser Frage kann die Iterationstheorie (*theory of runs*) herangezogen werden. Während die meisten Stichprobenfunktionen nur die Anzahl der Stichprobenelemente berücksichtigen, erlaubt die Iterationstheorie auch aus der Reihenfolge, mit der die einzelnen Stichprobenelemente auftreten, Schlüsse zu ziehen.

Unter einer Iteration (*run*) wird eine Folge identischer Stichprobenelemente verstanden, der Stichprobenelemente eines anderen Typs vorangehen oder folgen. Formal kann eine Iteration folgendermaßen definiert werden:

Gegeben sei eine Stichprobe mit den Werten $x_1, x_2, \dots, x_n$. Die Folge $x_j, \dots, x_{j+m}$ mit $j=1, 2, \dots, n$ und $m = 0, 1, \dots, n-j$ wird eine Iteration genannt, wenn

$$x_{j-1} \neq x_j = \dots = x_{j+m} \neq x_{j+m+1}$$

ist. Da $m = 0$ sein kann, kann eine Iteration auch aus einem einzigen Element bestehen.

Für die Anwendung der Iterationstheorie ist vor allem der Fall von Bedeutung, daß die untersuchte Stichprobe nur aus zwei Typen von Elementen besteht. Da wir im Kapitel 12.5 auch binäre Iterationen untersuchen werden und da dieser Typ von Iterationen zudem mathematisch weit einfacher zu behandeln ist, betrachten wir zunächst den Fall, daß die Stichprobe lediglich aus einer Folge der Elemente a_1 und a_2 besteht.

Gegeben sei z.B. die Stichprobe

$$\underline{a_1 a_2 a_2 a_2 a_1 a_1}.$$

Diese Folge von $n = 7$ Elementen besteht aus $n_1 = 3$ a_1 -Elementen und aus $n_2 = 4$ a_2 -Elementen. Sie enthält eine a_1 -Iteration der Länge 1, eine a_1 -Iteration der Länge 2 und eine a_2 -Iteration der Länge 4. Die Gesamtzahl der a_1 -Iterationen beträgt $r_1 = 2$, die der a_2 -Iterationen $r_2 = 1$. Insgesamt befinden sich in der Stichprobe $r_1 + r_2 = r = 3$ Iterationen.

Bei der Ermittlung der Verteilung der Iterationszahlen sind zwei Fälle zu unterscheiden:

(1) Am Anfang und am Ende der Stichprobe findet sich der gleiche Iterationstyp. In diesem Fall ist $|r_1 - r_2| = 1$, und die Stichprobe besteht entweder aus $r_1 = k$ a_1 -Iterationen und $r_2 = k + 1$ a_2 -Iterationen oder aus $r_1 = k + 1$ a_1 -Iterationen und $r_2 = k$ a_2 -Iterationen. Die Gesamtzahl der Iterationen beträgt dann $r = k + 1 + k = 2k + 1$.

(2) Am Anfang und am Ende der Stichprobe finden sich unterschiedliche Iterationstypen. In diesem Fall ist $r_1 = r_2 = k$, und die Gesamtzahl der Iterationen ist $r = 2k$.

Wir suchen nun die Wahrscheinlichkeitsfunktionen

$$P(r = 2k) \quad \text{und} \quad P(r = 2k + 1).$$

Um $P(r = 2k)$ zu finden, überlegt man sich zuerst, wieviele Möglichkeiten es gibt, n_1 a_1 -Elemente in k Iterationen zu unterteilen. Unter $n_1 - 1$ möglichen Zwischenräumen müssen $k - 1$ ausgewählt werden, und hierfür gibt es $\binom{n_1 - 1}{k - 1}$ Möglichkeiten. n_2 a_2 -Elemente können dementsprechend auf $\binom{n_2 - 1}{k - 1}$ verschiedene Arten unterteilt werden, um k Iterationen zu erhalten. Berücksichtigt man noch, daß die Stichprobe entweder mit einer a_1 -Iteration beginnt und mit einer a_2 -Iteration aufhört oder mit einer a_2 -Iteration beginnt und mit einer a_1 -Iteration aufhört, dann gibt es insgesamt

$$2 \binom{n_1 - 1}{k - 1} \binom{n_2 - 1}{k - 1}$$

Möglichkeiten, $2k$ Iterationen zu erhalten.

Die Zahl der möglichen Anordnungen einer Folge der Länge $n_1 + n_2 = n$ beträgt insgesamt $\binom{n}{n_1} = \binom{n}{n_2}$. Da alle Permutationen die gleiche Auftretenswahrscheinlichkeit haben, ist die Wahrscheinlichkeit einer beliebigen Anordnung

$$\frac{1}{\binom{n}{n_1}} = \frac{1}{\binom{n}{n_2}}. \quad (12.65)$$

Die gesuchte Wahrscheinlichkeit von $r = 2k$ Iterationen beträgt somit

$$P(r = 2k) = \frac{2 \binom{n_1 - 1}{k - 1} \binom{n_2 - 1}{k - 1}}{\binom{n}{n_1}}. \quad (12.66)$$

Ist $r_1 = k + 1$ und $r_2 = k$, dann gibt es $\binom{n_1 - 1}{k}$ Möglichkeiten, n_1 a_1 -Elemente in $k + 1$ Iterationen zu unterteilen, $\binom{n_2 - 1}{k - 1}$ Möglichkeiten, n_2 a_2 -Elemente in k Iterationen zu unterteilen und insgesamt

$$\binom{n_1 - 1}{k} \binom{n_2 - 1}{k - 1}$$

Möglichkeiten, $2k + 1$ Iterationen zu erhalten. Ist $r_1 = k$ und $r_2 = k + 1$, gibt es dementsprechend

$$\binom{n_2 - 1}{k} \binom{n_1 - 1}{k - 1}$$

Möglichkeiten. Berücksichtigt man noch (12.65), dann erhält man

$$P(r = 2k + 1) = \frac{\binom{n_1 - 1}{k} \binom{n_2 - 1}{k - 1} + \binom{n_2 - 1}{k} \binom{n_1 - 1}{k - 1}}{\binom{n}{n_1}} \quad (12.67)$$

Erweitert man (12.67) mit $(n_1 - k)(n_2 - k)$, ergibt sich nach einigen Umformungen

$$P(r = 2k + 1) = \frac{2 \binom{n_1 - 1}{k - 1} \binom{n_2 - 1}{k - 1}}{\binom{n}{n_1}} \frac{n - 2k}{2k} \quad (12.68)$$

und damit die Rekursionsbeziehung

$$P(r = 2k + 1) = \frac{n - 2k}{2k} P(r = 2k). \quad (12.69)$$

Die Verteilung der Anzahl der a_1 -Iterationen bzw. der a_2 -Iterationen ohne Berücksichtigung, ob $r_1 = r_2$ oder $|r_1 - r_2| = 1$ ist, erhalten wir auf ähnliche Weise wie (12.66) und (12.67). Die beiden Wahrscheinlichkeitsfunktionen lauten:

$$P(r_1) = \frac{\binom{n_1 - 1}{r_1 - 1} \binom{n_2 + 1}{r_1}}{\binom{n}{n_1}} \quad (12.70)$$

$$P(r_2) = \frac{\binom{n_1 + 1}{r_2} \binom{n_2 - 1}{r_2 - 1}}{\binom{n}{n_1}} \quad (12.71)$$

Die kumulative Wahrscheinlichkeitsverteilung von r , d.h. $F(r)$, ist für $n_1, n_2 \leq 20$ von SWED/EISENHART (1943) tabelliert worden (abgedruckt z.B. bei OWEN 1962:377ff). Da die Berechnung von $F(r)$ sehr umständlich ist, benutzt man bei $n_1, n_2 > 20$ meist die Approximation durch die Normalverteilung. Hierzu benötigt man den Erwartungswert und die Varianz von r .

Um $E(r)$ und $V(r)$ zu ermitteln, berechnen wir zuerst die Erwartungswerte von r_1 und r_2 . Diese erhält man relativ einfach über die faktoriellen Momente (vgl. WILKS 1962:148).¹¹ Für r_1 lautet das g -te faktorielle Moment

$$E(r_1[g]) = \sum_{r_1=g}^{n_1} r_1[g] \frac{\binom{n_1 - 1}{r_1 - 1} \binom{n_2 + 1}{r_1}}{\binom{n}{n_1}},$$

wobei $r_1[g] = r_1(r_1 - 1) \dots (r_1 - g + 1)$ ist.

Unter Benutzung der Identität

$$\sum_{r_1=g}^{n_1} \binom{n_1 - 1}{r_1 - 1} \binom{n_2 + 1 - g}{r_1 - g} = \binom{n - g}{n_1 - g}$$

erhält man

$$E(r_1[g]) = \frac{(n_2 + 1)[g] \binom{n - g}{n_1 - g}}{\binom{n}{n_1}} \quad (12.72)$$

Auf analoge Weise ermittelt man

$$E(r_2[g]) = \frac{(n_1 + 1)[g] \binom{n - g}{n_2 - g}}{\binom{n}{n_1}} \quad (12.73)$$

Aus (12.72) und (12.73) folgt

$$E(r_1) = \frac{n_1(n_2 + 1)}{n} \quad (12.74)$$

$$E(r_2) = \frac{n_2(n_1 + 1)}{n} \quad (12.75)$$

$$\begin{aligned} V(r_1) &= E(r_1^{[2]}) + E(r_1) - (E(r_1))^2 \\ &= \frac{n_1(n_1 - 1)n_2(n_2 + 1)}{n^2(n - 1)} \end{aligned} \quad (12.76)$$

$$\begin{aligned} V(r_2) &= E(r_2^{[2]}) + E(r_2) - (E(r_2))^2 \\ &= \frac{n_2(n_2 - 1)n_1(n_1 + 1)}{n^2(n - 1)} \end{aligned} \quad (12.77)$$

Aus (12.74) und (12.75) erhalten wir

$$E(r) = E(r_1) + E(r_2) = \frac{2n_1n_2 + n}{n} \quad (12.78)$$

Aus der Formel für die Kovarianz (vgl. MOOD 1940)

$$\text{COV}(r_1, r_2) = \frac{n_1(n_1 - 1)n_2(n_2 - 1)}{n^2(n - 1)} \quad (12.79)$$

und aus (12.76) und (12.77) folgt schließlich

$$\begin{aligned} V(r) &= V(r_1) + V(r_2) + 2 \text{COV}(r_1, r_2) \\ &= \frac{2n_1n_2(2n_1n_2 - n)}{n^2(n - 1)} \end{aligned} \quad (12.80)$$

Die gesuchte Teststatistik lautet jetzt

$$\hat{z} = \frac{\hat{F} - E(\hat{F})}{\sqrt{V(\hat{F})}} = \frac{\hat{F} - \frac{2n_1n_2 + n}{n}}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n)}{n^2(n - 1)}}} \quad (12.81)$$

oder nach einigen Umformungen

$$\hat{z} = \frac{n(\hat{F} - 1) - 2n_1n_2}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n)}{n - 1}}} \quad (12.82)$$

Diese Approximation ist bereits bei $n_1, n_2 > 10$ relativ gut und kann durch eine Kontinuitätskorrektur noch verbessert werden.¹²

Die abgeleiteten Verteilungen lassen sich für den Fall von k verschiedenen Typen von Elementen verallgemeinern (vgl. z.B. MOOD 1940; DAVID/BARTON 1962:119ff). Die exakten Verteilungen sind jedoch, wie MOOD (1940) selbst feststellt, z.T. für die Praxis wenig geeignet. Wir benutzen deshalb zum Testen der Gesamtzahl der Iterationen von k verschiedenen Typen von Elementen wiederum die Normalverteilung, die auch im Fall von multiplen Iterationen bereits bei relativ kleinen Stichproben eine gute Approximation ermöglicht (vgl. BARTON/DAVID 1957:173).

Ersetzen wir in (12.74) und (12.76) n_1 durch n_i und n_2 durch $n - n_i$, erhalten wir den Erwartungswert und die Varianz der Anzahl der Iterationen eines bestimmten Typs a_i im Fall von k verschiedenen Typen von Elementen:

$$E(r_i) = \frac{n_i(n - n_i + 1)}{n} \quad (12.83)$$

$$V(r_i) = \frac{n_i(n_i - 1)(n - n_i)(n - n_i + 1)}{n^2(n - 1)} \quad (12.84)$$

Ersetzen wir in (12.79) n_1 durch n_i und n_2 durch n_j , ergibt sich die Formel für die Kovarianz k verschiedener Iterationszahlen:

$$\text{COV}(r_i, r_j) = \frac{n_i(n_i - 1)n_j(n_j - 1)}{n^2(n - 1)} \quad (12.85)$$

Anhand von (12.83), (12.84) und (12.85) kann nun der Erwartungswert und die Varianz der Gesamtzahl der Iterationen ermittelt werden:

$$E(r) = \sum_{i=1}^k E(r_i) = \sum_{i=1}^k \frac{n_i(n - n_i + 1)}{n} = n + 1 - \frac{\sum_{i=1}^k n_i^2}{n} \quad (12.86)$$

$$V(r) = \sum_{i=1}^k V(r_i) + \sum_{i=1}^k \sum_{\substack{j=1 \\ i \neq j}}^k \text{COV}(r_i, r_j)$$

Um die Doppelsumme besser auswerten zu können, erweitern wir den Summationsbereich auch auf den Fall $i=j$. Die dadurch erhaltene zusätzliche Summe ziehen wir wieder ab:

$$\begin{aligned} V(r) &= \sum_{i=1}^k \frac{n_i(n_i - 1)(n - n_i)(n - n_i + 1)}{n^2(n - 1)} \\ &+ \sum_{i=1}^k \sum_{j=1}^k \frac{n_i(n_i - 1)n_j(n_j - 1)}{n^2(n - 1)} \\ &- \sum_{i=1}^k \frac{n_i(n_i - 1)n_i(n_i - 1)}{n^2(n - 1)} \end{aligned}$$

Nach einigen Umformungen ergibt sich schließlich

$$V(r) = \frac{\sum_{i=1}^k n_i^2 \left[\sum_{i=1}^k n_i^2 + n(n + 1) \right] - 2n \sum_{i=1}^k n_i^3 - n^3}{n^2(n - 1)} \quad (12.87)$$

Transformiert man jetzt $\hat{Z} = (\hat{A} - E(\hat{A})) / \sqrt{V(\hat{A})}$, erhält man die gesuchte Teststatistik:

$$\hat{Z} = \frac{\sum n_i^2 - n(n - \hat{r})}{\sqrt{\frac{\sum n_i^2 \left[\sum n_i^2 + n(n + 1) \right] - 2n \sum n_i^3 - n^3}{n - 1}}} \quad (12.88)$$

mit $i = 1, \dots, k$

Mit der Teststatistik (12.88) kann nun geprüft werden, ob die beobachtete Gesamtzahl der Iterationen der Verstypen V_8 bis V_{12} von dem unter der Nullhypothese zu erwartenden Wert abweicht. Dazu berechnen wir zuerst die Erwartungswerte der einzelnen Verstypen. Die Anzahl der Iterationen des Typs V_9 z.B. beträgt 7 (vgl. Tab. 12.21). Da dieser Verstyp 16 mal im "Erlkönig" vorkommt, erhalten wir mit Formel (12.83)

$$E(\hat{A}_{V_9}) = \frac{16(16 + 1)}{32} = 8,50.$$

Die beobachteten und die erwarteten Iterationszahlen der einzelnen Verstypen sind in der Tabelle 12.23 aufgeführt.

Tab. 12.23: Iterationswerte der Verstypen aus Tabelle 12.21

Silben	n_i	$\hat{A}_i$	$E(\hat{A}_i)$	K_i
8	2	1	1,94	0,52
9	16	7	8,50	0,82
10	7	6	5,69	1,05
11	6	3	5,06	0,59
12	1	1	1,00	1,00
	32	18	22,19	0,81

Wir testen nun, ob sich die Gesamtzahl der Iterationen, d.h. $\hat{A} = 18$, signifikant vom Erwartungswert $E(\hat{A}) = 22,1875$ unterscheidet. Mit Hilfe der Formel (12.87) erhalten wir für die Varianz von r den Wert

$$V(\hat{A}) = \frac{346(346 + 32(33)) - 2(32)4664 - 32^3}{32^2(31)} = 4,8459.$$

Führen wir noch eine Kontinuitätskorrektur durch - da wir zweiseitig testen, vermindern wir $|\hat{A} - E(\hat{A})|$ um 0,5 - ergibt der Normalverteilungstest

$$Z = \frac{|\hat{A} - E(\hat{A})| - 0,5}{\sqrt{V(\hat{A})}} = \frac{|18 - 22,1875| - 0,5}{\sqrt{4,8459}} = 1,6751.$$

Da bei $P = 0,05$ der kritische Wert 1,96 beträgt, muß die Nullhypothese beibehalten werden. Es besteht somit im "Erlkönig" keine Tendenz, Verse mit gleicher Silbenzahl aufeinanderfolgen zu lassen.

Als Maß für die Stärke einer Tendenz zur Iteration kann nach einem Vorschlag von WORONCZAK (1961) der Quotient aus der beobachteten und der erwarteten Gesamtzahl der Iterationen verwendet werden:

$$K = \frac{\hat{A}}{E(\hat{A})} \quad (12.89)$$

Ebenso kann ein Index für die Stärke der Iteration eines bestimmten Typs gebildet werden:

$$K_i = \frac{\hat{A}_i}{E(\hat{A}_i)} \quad (12.90)$$

Bei einer zufälligen Anordnung der Stichprobenelemente ist der Iterationsindex gleich oder nahe 1. Liegt ein 'Klumpungseffekt' vor, ist K kleiner als 1, bei einer Tendenz zum regelmäßigen Wechsel dagegen größer als 1.

Die Werte für K und K_i sind in der letzten Spalte der Tabelle 12.23 aufgeführt. Wie zu erwarten war, weist der Gesamtindex und die Mehrzahl der Einzelindizes auf einen Klumpungseffekt hin.

Ebenso wie die Gesamtzahl der Iterationen kann auch die Anzahl der Iterationen eines bestimmten Typs auf Signifikanz getestet werden (vgl. MOOD 1940; DAVID/BARTON 1962:119ff). Weiterhin kann auch die längste Iteration als Testkriterium für die Zufälligkeit einer Folge benutzt werden. Kennt man die Auftretenswahrscheinlichkeit der Elemente der Folge in der Grundgesamtheit oder kann man z.B. annehmen, daß unter der Nullhypothese $n_1/n = p = n_2/n = q$ gilt, dann kann mit der Iterationstheorie nicht nur die Anzahl der Iterationen bei vorgegebenem n_i getestet werden, sondern es kann auch nach der Wahrscheinlichkeit gefragt werden, daß z.B. n_1 a_1 -Elemente und n_2 a_2 -Elemente in der Folge vorkommen und daß diese in einer bestimmten Weise angeordnet sind. Die Formeln für solche zusammengesetzten Wahrscheinlichkeiten erhält man einfach dadurch, daß man die entsprechenden Formeln - in unserem Fall die Formeln (12.66), (12.67), (12.70) und (12.71) - mit der binomialen Wahrscheinlichkeit

$$\binom{n}{n_1} p^{n_1} q^{n_2}$$

multipliziert.

Mit der Iterationstheorie lassen sich eine Vielzahl interessanter Fragestellungen in der Metrik untersuchen. Eine weitere Anwendungsmöglichkeit dieses Modells soll im folgenden Kapitel aufgezeigt werden. Neben der Iterationstheorie gibt es noch einige weitere statistische Modelle, um die Zufälligkeit der Anordnung von Stichprobenelementen zu prüfen. Eine Reihe interessanter Untersuchungsmöglichkeiten eröffnet z.B. die *theory of random clumping* (vgl. hierzu ROACH 1968, GREGER 1972), die m.W. bisher noch nicht in der Metrik verwendet worden ist. Ein weiteres Modell ist die Korrelationsrechnung. Auf dieses Modell werde ich in Kapitel 12.5.2 näher eingehen.

12.5. Analyse der metrischen Bindung

12.5.1. In den Kapiteln 4 und 5 hatte ich als Definitionsmerkmal des Prädikats 'metrisch' u.a. die objektiv feststellbare Periodizität von Silbenfolgen genannt. In diesem Kapitel soll nun gezeigt werden, wie mit Hilfe statistischer Methoden festgestellt werden kann, ob und in welchem Maße eine Silbenfolge als periodisch anzusehen ist.

Die Frage, ob eine Folge als periodisch anzusehen ist, läßt sich wiederum mit Hilfe der Iterationstheorie beantworten.

Die Tabelle 12.24 zeigt die beobachtete und erwartete Anzahl der Iterationen von markierten und unmarkierten Silben im "Erlkönig" und in "Der Totentanz". Zur Berechnung der Erwartungswerte wurden die Formeln (12.74) und (12.75) benutzt.

Tab. 12.24: Beobachtete und erwartete Anzahl der Iterationen von markierten (M) und unmarkierten Silben (U) im "Erlkönig" und in "Der Totentanz"

Silben	Erlkönig			Totentanz		
	n_i	$\hat{A}_i$	$E(\hat{A}_i)$	n_i	$\hat{A}_i$	$E(\hat{A}_i)$
M	128	126	75,22	175	175	113,48
U	180	126	75,39	320	176	113,78
	308	252	150,61	495	351	227,26

Da aufgrund des Metrums Iterationen von mehr als einer markierten Silbe im "Erlkönig" nur 'ausnahmsweise' am Anfang der Verse 22 und 28 und in "Der Totentanz" überhaupt nicht vorkommen - in "Der Totentanz" ist deshalb die Anzahl der Iterationen von markierten Silben gleich der Anzahl der markierten Silben - ist in beiden Gedichten die erwartete Anzahl der Iterationen deutlich geringer als die beobachtete Anzahl.

Um die Gesamtzahl der Iterationen auf Signifikanz testen zu können, benötigen wir noch die Varianzen. Mit Hilfe der Formel (12.80) erhalten wir für den "Erlkönig" $V(\hat{A}) = 72,42$ und für "Der Totentanz" $V(\hat{A}) = 103,18$. Der Normalverteilungstest ergibt die beiden folgenden Werte:

$$\text{"Erlkönig":} \quad \hat{z} = \frac{252 - 150,61}{8,51} = 11,91$$

$$\text{"Der Totentanz":} \quad \hat{z} = \frac{351 - 227,26}{10,16} = 12,18$$

Der $\hat{z}$ -Wert ist in beiden Fällen hochsignifikant. Die Hypothese einer Tendenz zur regelmäßigen Alternanz von markierten und unmarkierten Silben kann deshalb für beide Gedichte als bestätigt angesehen werden.

Mit Hilfe der Normalverteilung kann auch getestet werden, ob zwischen den beiden Gedichten ein Unterschied besteht. Die Teststatistik lautet

$$\hat{z} = \frac{\hat{A}_A - \hat{A}_B - (E(\hat{A}_A) - E(\hat{A}_B))}{\sqrt{V(\hat{A}_A) + V(\hat{A}_B)}} \quad (12.91)$$

Mit diesem Testkriterium wird der Unterschied zweier Iterationszahlen nicht direkt geprüft, sondern indirekt über den Unterschied der Abweichungen beider Iterationszahlen von ihren Erwartungswerten. Ein direkter Vergleich ist lediglich dann möglich, wenn in bezug auf die Anzahl der Elemente a_1 und a_2 in den zu vergleichenden Texten $n_{1A} \cdot n_{2A} = n_{1B} \cdot n_{2B}$ gilt. In diesem Fall ist $E(\hat{A}_A) = E(\hat{A}_B)$ und (12.91) wird zu

$$\hat{z} = \frac{\hat{A}_A - \hat{A}_B}{\sqrt{V(\hat{A}_A) + V(\hat{A}_B)}} \quad (12.92)$$

Der Vergleich zwischen dem "Erlkönig" und "Der Totentanz" ergibt

$$\hat{z} = \frac{252 - 351 - (150,61 - 227,26)}{\sqrt{72,42 + 103,18}} = 1,69.$$

Bei $P = 0,05$ beträgt der kritische Wert 1,96. Der Unterschied zwischen beobachteter und erwarteter Anzahl der Iterationen ist somit in beiden Texten als gleich anzusehen.

Berechnet man den Iterationsindex (vgl. Formel 12.89), erhalten wir für den "Erlkönig" $K_E = 252/150,61 = 1,67$ und für "Der Totentanz" $K_T = 351/227,26 = 1,54$. Der Index ist in beiden Fällen erheblich größer als 1 und zeigt damit eine deutlich ausgeprägte Tendenz zur Alternanz an. Wegen des höheren Wertes für den "Erlkönig" könnte man meinen, daß diese Tendenz im "Erlkönig" stärker ausgeprägt ist. An dieser Stelle zeigt sich ein Nachteil des von WORONCZAK vorgeschlagenen Iterationsindexes. Dieser Index ist nämlich nicht vom Stichprobenumfang unabhängig. Je größer die Stichprobe ist, desto deutlicher weicht K auch bei gleich gebauten Folgen von 1 ab. Ein direkter Vergleich von Indexwerten ist nur dann möglich, wenn die Erwartungswerte der Iterationen in den zu vergleichenden Stichproben übereinstimmen.

Es stellt sich deswegen die Frage, ob es nicht sinnvoller ist, den Index dadurch zu normieren, daß man die beobachtete Anzahl der Iterationen auf das jeweilige theoretisch mögliche Minimum bzw. Maximum bezieht. Auf diese Weise ist es möglich, eine Normierung auf das Einheitsintervall (0; 1) vorzunehmen und damit Indexwerte direkt vergleichbar zu machen. Im Fall von binären Iterationen ist das theoretische Minimum bzw. Maximum ohne Schwierigkeiten zu ermitteln. Ist $n_1 = n_2$, ist $\max(r) = 2n_1 = 2n_2 = n$. Ist $n_1 \neq n_2$, ist $\max(r) = 2 \cdot \min(n_1, n_2) + 1$. In bezug auf das Minimum gilt $\min(r) = 2$. Für binäre Iterationen können wir den gesuchten Index nunmehr definieren als

$$It = \frac{\hat{A} - 2}{\max(r) - 2} \quad (12.93)$$

Dieser Index ist gleich 1, wenn $\hat{A} = \max(r)$ und gleich 0, wenn $\hat{A} = \min(r) = 2$ ist.

In "Der Totentanz" ist $\min(n_1, n_2) = 175$ und damit $\max(r) = 2(175) + 1 = 351$. Wir erhalten somit

$$It = \frac{351 - 2}{351 - 2} = 1.$$

Für den "Erlkönig" beträgt It wegen der beiden 'Ausnahmen' in den Versen 22 und 28 dagegen nur 0,9804.

Bisher haben wir lediglich die Abfolge der Silben in metrischen Texten untersucht. Die Tabelle 12.25 zeigt die beobachtete und erwartete Anzahl der Iterationen in einem Vers- und in einem Prosatext. Bei den beiden Stichproben handelt es sich um die Textanfänge von VERGILs "Aeneis" und CAESARS "BELLUM GALLICUM". Unter einer markierten Silbe ist dabei entweder eine natur- oder eine positions-lange Silbe zu verstehen.

Tab. 12.25: Beobachtete und erwartete Anzahl der Iterationen in der "Aeneis" und im "Bellum Gallicum"

Silben- typ	Aeneis			Bellum Gallicum		
	n_i	$\hat{A}_i$	$E(\hat{A}_i)$	n_i	$\hat{A}_i$	$E(\hat{A}_i)$
M	45	30	17,50	48	15	14,82
U	27	14	17,25	20	13	14,41
	72	44	34,75	68	28	29,24

Der Normalverteilungstest ergibt folgende Werte:

"Aeneis": $\hat{z} = \frac{44 - 34,75}{\sqrt{15,57}} = 2,34$

"Bellum Gallicum": $\hat{z} = \frac{28 - 29,24}{\sqrt{11,48}} = -0,34$

Vergleichen wir die beiden Texte untereinander, erhalten wir

$$\hat{z} = \frac{44 - 28 - (34,75 - 29,24)}{\sqrt{15,57 + 11,48}} = 2,02.$$

Die $\hat{z}$ -Werte zeigen deutlich den Unterschied zwischen dem Vers- und dem Prosatext. Im Prosatext stimmt die erwartete und die beobachtete Anzahl der Iterationen weitgehend überein. Die Anordnung der Silben kann in diesem Text als zufällig angesehen werden. Anhand des Vers-textes zeigt sich dagegen deutlich, daß das Metrum die Wahlmöglichkeit einengt und die sequentielle Wahrscheinlichkeitsstruktur der Sprache verstärkt. Informationstheoretisch kann diese Tendenz als Tendenz zur Verstärkung der normalsprachlichen Redundanz interpretiert werden. Diese zusätzliche Redundanz metrische Texte dürfte auch der Grund dafür sein, daß metrische Texte in der Regel leichter memorisierbar sind als Prosatexte.¹³

12.5.2. Das vor allem von W. FUCKS zur Untersuchung von Texten verwendete Modell der Korrelationsrechnung (vgl. z.B. FUCKS 1955; 1968) erlaubt nicht nur die Zufälligkeit der Anordnung der Elemente einer Folge zu überprüfen, sondern darüber hinaus auch die Stärke und Reichweite von Periodizitäten zu messen. Mit diesem Modell lassen sich sowohl spezielle Fragen der Metrik als auch allgemeine textwissenschaftliche Phänomene untersuchen. So hat FUCKS anhand eines umfangreichen Korpus z.B. festgestellt, daß die Satz-längen aufeinanderfolgender Sätze eines Textes nicht unabhängig voneinander sind, sondern daß lange Sätze eher zusammen mit langen Sätzen auftreten und kurze Sätze eher mit kurzen Sätzen (vgl. FUCKS 1968; 1971).

Um die Korrelationsrechnung anwenden zu können, betrachten wir den "Erlkönig" als eine aus markierten und unmarkierten Silben bestehende Folge, deren Anfang und Ende zyklisch miteinander verbunden ist. Wir ermitteln nun, wie oft auf eine markierte Silbe als nächste eine markierte Silbe folgt, auf eine unmarkierte eine unmarkierte, auf eine markierte eine unmarkierte und auf eine unmarkierte eine markierte. Die Kombination markiert - markiert kommt zweimal vor, unmarkiert - unmarkiert dagegen in 54 Fällen: $MM_1 = 2$, $UU_1 = 54$. Die Kombination markiert - unmarkiert und unmarkiert - markiert finden wir in 126 Fällen: $MU_1 = 126$, $UM_1 = 126$. Der Index 1 soll anzeigen, daß es sich um direkt benachbarte Silben handelt. Jetzt denken wir uns jede Silbe jeweils mit der übernächsten kombiniert. Daß es sich um die übernächste Silbe handelt, zeigt der Index 2 an: $MM_2 = 73$, $UU_2 = 125$, $MU_2 = 55$, $UM_2 = 55$. Um den "Erlkönig" mit den von FLOCKS für andere Werke ermittelten Werten vergleichen zu können, gehen wir bei unserer Berechnung bis zum Abstand 16. Dabei berücksichtigen wir auch den theoretischen Fall, daß der Abstand 0 beträgt. Hierfür erhalten wir $MM_0 = 128$, $UU_0 = 180$. Dies bedeutet, daß der "Erlkönig" aus 308 Silben besteht, von denen 128 markiert und 180 unmarkiert sind.

Die Ergebnisse der Auszählung schreiben wir in Matrizen folgender Form:

$$\begin{bmatrix} MM & MU \\ UM & UU \end{bmatrix}$$

Es ergeben sich insgesamt folgende Werte:

$$\begin{bmatrix} 128 & 0 \\ 0 & 180 \end{bmatrix}_0 \quad \begin{bmatrix} 2 & 126 \\ 126 & 54 \end{bmatrix}_1 \quad \begin{bmatrix} 73 & 55 \\ 55 & 125 \end{bmatrix}_2$$

$$\begin{bmatrix} 52 & 76 \\ 76 & 104 \end{bmatrix}_3 \quad \begin{bmatrix} 43 & 85 \\ 85 & 95 \end{bmatrix}_4 \quad \begin{bmatrix} 53 & 75 \\ 75 & 105 \end{bmatrix}_5$$

$$\begin{bmatrix} 40 & 98 \\ 88 & 92 \end{bmatrix}_6 \quad \begin{bmatrix} 63 & 65 \\ 65 & 115 \end{bmatrix}_7 \quad \begin{bmatrix} 46 & 82 \\ 82 & 98 \end{bmatrix}_8$$

$$\begin{bmatrix} 56 & 72 \\ 72 & 108 \end{bmatrix}_9 \quad \begin{bmatrix} 54 & 74 \\ 74 & 106 \end{bmatrix}_{10} \quad \begin{bmatrix} 55 & 73 \\ 73 & 107 \end{bmatrix}_{11}$$

$$\begin{bmatrix} 58 & 70 \\ 70 & 110 \end{bmatrix}_{12} \quad \begin{bmatrix} 45 & 83 \\ 83 & 97 \end{bmatrix}_{13} \quad \begin{bmatrix} 61 & 67 \\ 67 & 113 \end{bmatrix}_{14}$$

$$\begin{bmatrix} 46 & 82 \\ 82 & 98 \end{bmatrix}_{15} \quad \begin{bmatrix} 60 & 68 \\ 68 & 112 \end{bmatrix}_{16}$$

In Kapitel 12.5.1 hatten wir als Maß für den Zusammenhang zweier nominalskaliert dichter Variablen den Kontingenzkoeffizienten ϕ benutzt. Mit Hilfe dieses Koeffizienten läßt sich auch die Verbundenheit zwischen den Elementen ein und derselben Menge, d.h. deren Autokontingenz (Autokorrelation) messen. In der Notation dieses Kapitels lautet ϕ folgendermaßen:

$$\phi = \frac{MM(UU) - MU(UM)}{\sqrt{(MM + MU)(UM + UU)(MM + UM)(MU + UU)}}$$

Wir berechnen ϕ nun für die insgesamt 17 Matrizen. Für die Matrix 1 z.B. ergibt sich der Wert

$$\phi_1 = \frac{2(54) - 126(126)}{\sqrt{(2 + 126)(2 + 126)(126 + 54)(126 + 54)}} = -0,684.$$

Dieser Wert zeigt deutlich die Abhängigkeit zwischen direkt aufeinanderfolgenden Silben im "Erlkönig". Das negative Vorzeichen zeigt an, daß Silben ungleichen Typs bevorzugt aufeinanderfolgen.

Die Werte für die Kontingenzkoeffizienten sind in der Tabelle 12.26 aufgeführt.

Tab. 12.26: Kontingenzkoeffizienten für den Grad der metrischen Bindungen im "Erlkönig"

Matrix	φ	Matrix	φ
0	1,0000	9	0,0375
1	-0,6844	10	0,0108
2	0,2648	11	0,0241
3	-0,0160	12	0,0642
4	-0,1363	13	-0,1095
5	-0,0026	14	0,1043
6	-0,1764	15	-0,0962
7	0,1311	16	0,0910
8	-0,0962		

Mit Hilfe der Beziehung $\varphi^2 = \chi^2/n$ (vgl. Kap. 12.5.2) können wir testen, wann ein φ -Koeffizient als signifikant anzusehen ist. Wir bilden dazu ein Konfidenzintervall mit $P = 0,95$. Da bei $P = 0,95$ $\chi^2_1 = 3,84$ ist, erhalten wir

$$P(\varphi^2 < \frac{\chi^2}{n}) = P(|\varphi| < \sqrt{\frac{3,84}{308}}) = P(|\varphi| < 0,1117) = 0,95.$$

Der kritische Wert beträgt somit 0,1117. Dieser Wert wird lediglich von 6 der 17 Koeffizienten überschritten, und zwar von φ_0 , φ_1 , φ_2 , φ_4 , φ_6 und φ_7 . Sieht man einmal von dem Koeffizienten φ_0 ab, der notwendigerweise immer 1 ist, dann weist lediglich φ_1 auf eine stark ausgeprägte Autokontingenz hin.

Bei der Interpretation dieses Testresultats ist allerdings zu berücksichtigen, daß wir - sieht man von φ_0 ab - insgesamt 16 analoge Tests an ein und demselben Text durchgeführt haben. Diese 16 Tests können deswegen als eine "Familie" simultaner Tests aufgefaßt werden. Auf die Problematik der Abgrenzung von Familien simultaner Tests sei dabei an dieser Stelle lediglich hingewiesen (vgl. hierzu MILLER 1966:31-35). Folgt man der Auffassung, daß es sich um simultane Tests handelt, ergeben sich erhebliche Konsequenzen in bezug auf die *faktische* Irrtumswahrscheinlichkeit der durchgeführten Tests. Die faktische Irrtumswahrscheinlichkeit, d.h. die Wahrscheinlichkeit, die Nullhypothese abzulehnen, obwohl sie in Wirklichkeit zutrifft, ist nämlich im Fall simultaner Tests oft erheblich größer als das angesetzte nominelle Signifikanzniveau. Dies kann leicht am Beispiel *unabhängiger* simultaner Tests mit ein und derselben Nullhypothese gezeigt werden. In diesem Fall beträgt nämlich die Wahrscheinlichkeit, bei einem Signifikanzniveau von jeweils $P = 0,05$ unter z.B. 16 Testergebnissen mindestens ein signifikantes Ergebnis zu erhalten, obwohl in Wirklichkeit die Nullhypothese in allen Fällen gilt, nicht mehr 0,05, sondern

$$P(X \geq 1) = \sum_{x=1}^{16} \binom{16}{x} 0,05^x 0,95^{16-x} = 1 - P(X = 0) \\ = 1 - 0,95^{16} = 0,5599.$$

Eine mögliche Lösung des Problems besteht darin, die Irrtumswahrscheinlichkeiten für die einzelnen Hypothesen so zu wählen, daß die Summe der Irrtumswahrscheinlichkeiten der Einzelhypothesen die vorgegebene Irrtumswahrscheinlichkeit nicht überschreitet (vgl. KRAUTH/LIENERT 1973:40ff). Bezeichnet man die Wahrscheinlichkeit, daß mindestens ein signifikantes Ergebnis unter insgesamt r Ergebnissen auftritt, obwohl alle einzelnen Nullhypothesen zutreffen, mit P und die Irrtumswahrscheinlichkeit jeder einzelnen Hypothese mit P_i , dann muß also gelten

$$\sum_{i=1}^r P_i \leq P.$$

Fordert man für alle r simultanen Tests die gleiche Irrtumswahrscheinlichkeit P_i , dann ist diese Ungleichung erfüllt, wenn das vorgegebene Signifikanzniveau P für jeden einzelnen Test folgendermaßen adjustiert wird

$$P_i = P^* = P/r. \quad (12.94)$$

Diese Adjustierung berücksichtigt sowohl die aus der Abhängigkeit als auch aus der Häufung von Tests resultierenden Wirkungen (vgl. KRAUTH/LIENERT 1973:42). Für den vorliegenden Fall ergibt die Adjustierung $P^* = 0,05/16 = 0,0031$. Für jeden der 16 Tests ist somit ein Signifikanzniveau von $P^* = 0,0031$ anzusetzen.

Wir berechnen nun ein entsprechend adjustiertes Konfidenzintervall für ϕ . Für $P^* = 0,0031$ erhalten wir einen χ^2 -Wert von 8,7320 und eine kritische Schranke von $\phi = 0,1684$. Diese neue Schranke wird jetzt nur noch - läßt man ϕ_0 unberücksichtigt - von ϕ_1 , ϕ_2 und ϕ_6 überschritten.

Wie das Beispiel zeigt, führt die Adjustierung des Signifikanzniveaus bei einer größeren Anzahl von simultanen Tests zu einem relativ konservativen Testverfahren und damit zu einer Erhöhung des β -Fehlers. Es ist deshalb in jedem einzelnen Fall zu prüfen, ob eine Adjustierung dem zu untersuchenden Problem angemessen ist. Will man z.B. in einer explorativen Untersuchung lediglich eine Reihe interessanter Hypothesen testen und die Ergebnisse anschließend an neuem Material überprüfen, dann kann es durchaus sinnvoll sein, auf eine Adjustierung zu verzichten und so die Teststärke zu erhöhen. Will man jedoch aus den Alternativhypothesen weitreichende (praktische) Konsequenzen ziehen, dann dürfte auf jeden Fall ein simultanes Testverfahren bzw. eine Adjustierung des Signifikanzniveaus und damit ein erhöhter Schutz der Nullhypothese ratsam sein.

Stellt man die berechneten Kontingenzkoeffizienten graphisch dar, dann ergibt sich das in Abbildung 12.3 gezeigte Korrelogramm.

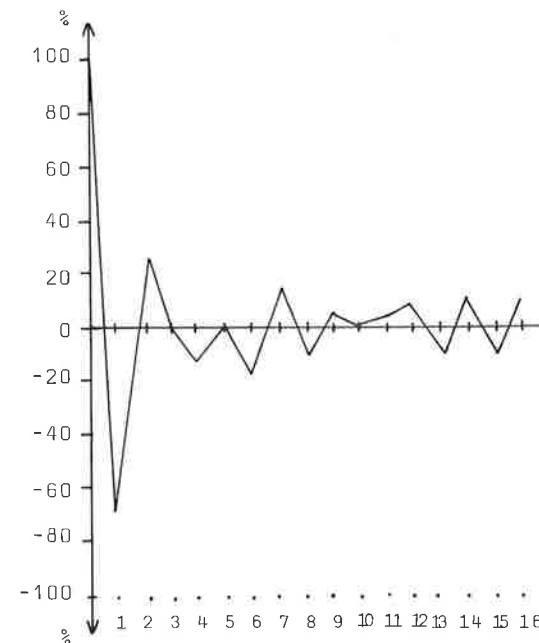


Abb. 12.3: Korrelogramm der metrischen Bindungen des "Erlkönig" (%)

Die positiven Korrelationskoeffizienten des Korrelogramms zeigen, daß gleichartige Markierungen (MM, UU) bevorzugt aufeinander folgen, negative Korrelationskoeffizienten zeigen, daß ungleiche Markierungen bevorzugt werden.

Zum Vergleich seien zwei aus FUCKS (1968:69) entnommene Korrelogramme reproduziert.

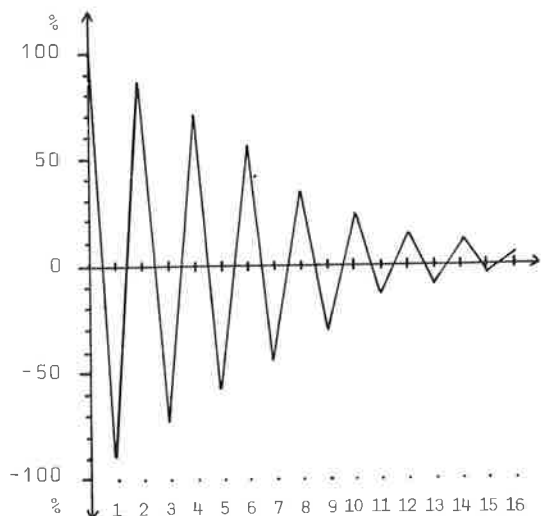


Abb. 12.4: GOETHE, Torquato Tasso, Jambus (u b)

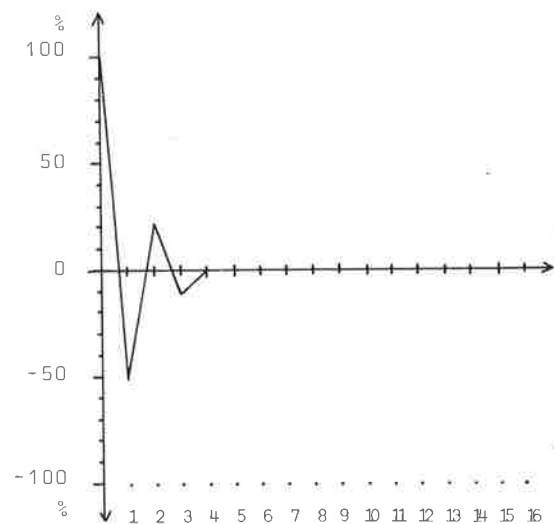


Abb. 12.5: BENN, Prosa und Szenen

Der Unterschied zwischen dem Bau des Volksliedverses des "Erlkönig" und den Jamben des "Torquato Tasso" tritt deutlich zutage. Die Kontingenzkoeffizienten des "Torquato Tasso" weisen wesentlich höhere Werte auf, die auch nur langsam abnehmen. Außerdem wechseln negative und positive Koeffizienten regelmäßig. Es zeigt sich wiederum, daß der Volksliedvers metrisch relativ frei ist. Das Korrelogramm der "Prosa und Szenen" BENNs, dessen Koeffizienten wesentlich niedriger sind und sehr schnell den Wert 0 erreichen, macht jedoch deutlich, daß immer noch ein großer Abstand zur Prosa besteht. Wie die weiteren Korrelogramme bei FUCHS (1968) zeigen, bestehen jedoch auch zwischen verschiedenen Prosatexten erhebliche Unterschiede. So finden sich z.B. in ADENAUERS "Erinnerungen" deutlich niedrigere Koeffizienten als in BENNs "Prosa und Szenen".

Vergleicht man die Korrelogramme 12.3, 12.4 und 12.5 mit Darstellungen sog. gedämpfter Schwingungen, die bei bestimmten physikalischen Vorgängen, wie z.B. beim Anschlagen einer Stimmgabel entstehen, zeigt sich eine relativ große Affinität.

Gedämpfte Schwingungen sind dadurch charakterisiert, daß ihre Amplitude, d.h. die Schwingungsweite wegen des stetigen Einwirkens der Widerstände auf die Schwingungsbewegung und wegen der Energieabstrahlung ständig abnimmt. Die Abnahme ist dabei nicht konstant, sondern erfolgt nach dem Gesetz einer Exponentialreihe (vgl. v. ESSEN 1961:53).

Gedämpfte Schwingungen lassen sich mit Hilfe verschiedener mathematischer Verfahren modellieren und meist durch relativ einfache Funktionen beschreiben. Verwendet man die bei v. ESSEN (1961) geschilderten Verfahren, ergibt sich für das Erlkönig-Korrelogramm aufgrund des nicht sehr regelmäßigen Wechsels zwischen positiven und negativen Koeffizienten und der dadurch gegebenen partiellen Aperiodizität keine befriedigende Approximation.

Eine sehr gute Approximation läßt sich jedoch für das absolut periodische Korrelogramm des "Torquato Tasso" erzielen. Da die Amplitude dieses Korrelogramms zumindest im Bereich der ersten 10 Koeffizienten nicht exponential, sondern linear abnimmt, gehen wir

bei der Suche nach einer geeigneten Funktion von einer trigonometrischen Funktion der allgemeinen Form

$$y = \cos(bx + a)$$

aus. Die Parameter dieser Funktion können wiederum mit der Methode der kleinsten Quadrate ermittelt werden (vgl. Kap. 12.2.2). Dazu berechnet man zuerst die Werte der Funktion $z = \cos^{-1}\varphi = \arccos \varphi$, d.h. man ermittelt die zu den Koeffizienten gehörenden Winkel. Anschließend setzt man $z = ax + b$ und schätzt die Parameter dieser Gleichung mit Hilfe der Formeln (12.7) und (12.11).

Bei der Regressionsrechnung lassen wir den Koeffizienten φ_0 , der immer 1 ist und deswegen keinen Informationswert hat, unberücksichtigt. Außerdem vernachlässigen wir die zur Charakterisierung der Stärke der metrischen Bindung nicht sehr bedeutsamen Koeffizienten φ_{11} bis φ_{16} . Da sich bei FÜCKS (1968) keine Angaben über die exakten Werte der Kontingenzkoeffizienten finden, legen wir bei der Regressionsrechnung die aus dem Korrelogramm des "Torquato Tasso" abgelesenen, approximativen Werte zugrunde.

Wir erhalten $b = 174,42$ und $a = -22,91$. Die gesuchte Funktion lautet somit

$$y = \cos(174,42x - 22,91). \quad (12.95)$$

Die beobachteten und theoretischen Werte der Kontingenzkoeffizienten sind in der Tabelle 12.27 aufgeführt. Es zeigt sich, daß die Funktion (12.95) zumindest bis zum Abstand von 10 Silben sehr genau den Verlauf der Koeffizienten des "Torquato Tasso" beschreibt.¹⁴

Tab. 12.27: Beobachtete (φ_b) und theoretische (φ_t) Werte der Kontingenzkoeffizienten des "Torquato Tasso".

Abstand	φ_b	φ_t	Abstand	φ_b	φ_t
1	-0,90	-0,88	6	0,55	0,55
2	0,85	0,83	7	-0,45	-0,46
3	-0,75	-0,77	8	0,35	0,38
4	0,70	0,70	9	-0,30	-0,29
5	-0,60	-0,63	10	0,25	0,20

Es soll jetzt außer Betracht gelassen werden, ob Verbundenheit zwischen gleichmarkierten oder ungleichmarkierten Silben besteht. Es soll nur noch untersucht werden, ob Verbundenheit besteht, in welchem Maße und wie weit sie reicht. Wir berücksichtigen deshalb lediglich die Absolutbeträge der ermittelten Korrelationskoeffizienten. Diese ergeben folgendes Korrelogramm:

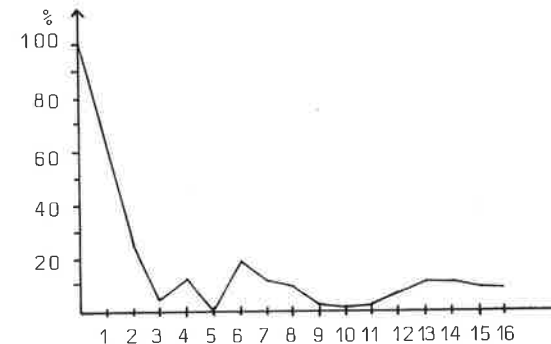


Abb. 12.6: Absolutbeträge der Korrelationskoeffizienten des "Erlkönig"

Geht man von dem Korrelogramm der absoluten Werte der ϕ -Koeffizienten aus, dann kann die Fläche zwischen dem Streckenzug und der Abszissenachse als Maß für die Stärke der metrischen Bindung angesehen werden. Die Fläche ist gleich der Summe der Kontingenzkoeffizienten. Bestände zwischen den Silben eines Textes eine vollkommene metrische Bindung, ergäbe sich jedesmal ein Koeffizient mit dem Wert 1 und, da wir 17 Werte berechnet haben, insgesamt der Wert 17. Entsprechend einem Vorschlag von FUCKS (1968:76) kann man nun das Verhältnis der errechneten zu der theoretisch möglichen Stärke als einen Index für die Stärke der metrischen Bindung (S_B) ansehen. Schreibt man diesen Index als

$$S_B = \frac{1}{n} \sum_{i=0}^n \phi_i \quad (12.96)$$

dann wird deutlich, daß dieser Index auch als arithmetisches Mittel der Kontingenzkoeffizienten interpretiert werden kann.

Für den "Erlkönig" erhalten wir

$$S_B = \frac{3,0402}{17} = 0,1788.$$

Ein Vergleich mit den von FUCKS (1968:73) für 15 Prosa- und 20 Vers-
texte ermittelten Werte zeigt, daß der "Erlkönig" sich zwar deutlich
von Prosatexten in bezug auf die Stärke seiner metrischen Bindung
unterscheidet - den höchsten Prosawert nimmt BENN mit etwa 11 % ein -,
aber innerhalb der Verstexte zu den Werken mit geringer metrischer
Bindung zu zählen ist.¹⁵

Mit der Korrelationsrechnung lassen sich eine Vielzahl textueller
Phänomene untersuchen. Dabei kann natürlich nicht nur der Zusammen-
hang qualitativer Merkmale, wie markiert und unmarkiert, sondern
auch die Autokorrelation zwischen quantitativen Merkmalen festge-
stellt werden. Das Vorgehen ist dabei, abgesehen von der Verwendung
anderer Korrelationskoeffizienten, prinzipiell gleich.

12.6. Die Verteilung der Anzahl der Silben pro Wort

Als eine weitere Möglichkeit zur Charakterisierung von Texten
kann die Anzahl der Silben pro Wort angesehen werden. Um die be-
kannten Schwierigkeiten bei der Definition des Wortes als seman-
tischer Einheit zu vermeiden, soll im folgenden unter einem Wort eine
nicht durch Satzzeichen unterbrochene Folge von Graphemen zwischen
zwei Leerstellen verstanden werden. Diese operationale Wortdefini-
tion hat noch den weiteren Vorteil, daß sie den Einsatz elektronischer
Datenverarbeitungsanlagen bei der Textverarbeitung ermöglicht (vgl.
KRALLMANN 1966:78f).

Als erstes soll nun die Anzahl der Silben pro Wort im "Erlkönig"
und in 10 Briefen GOETHEs verglichen werden. Dabei werde ich auf
eine umfassendere statistische Analyse verzichten und vor allem einige
allgemeine Fragen diskutieren. Bei den Briefen handelt es sich um
eine Zufallsauswahl (mit Hilfe von Zufallszahlen) aus den insgesamt
97 Briefen des Jahres 1782. Dieses Jahr habe ich deshalb gewählt, um
beim Vergleich mit dem "Erlkönig", der ebenfalls 1782 verfaßt worden
ist, den Einfluß der Variablen 'Diachronie' (vgl. Kap. 11.2.3) aus-
zuschalten.

Die Daten für den "Erlkönig" und für die "Briefe" finden sich in
der Tabelle 12.28. Es fällt sofort auf, daß die "Briefe" im Gegen-
satz zum "Erlkönig" Fünf- und Sechssilbler enthalten. Im Brief 644
findet sich sogar ein achtsilbiges Wort. Außerdem ist der Anteil
der einsilbigen Wörter in den "Briefen" wesentlich geringer als im
"Erlkönig". Der Unterschied zwischen den beiden Textsorten zeigt sich
auch beim Vergleich der Mittelwerte und der Varianzen. Insgesamt
gesehen sind Mittelwert und Varianz in den "Briefen" erheblich größer
als im "Erlkönig". Das gleiche gilt für den Variationskoeffizienten
v.¹⁶

Tab. 12.28: Verteilung der Anzahl der Silben pro Wort im "Erlikönig" und in 10 Briefen GOETHEs
aus dem Jahre 1782

Silben	Erlikönig	B r i e f N r.											
		612		589		596		591		667			
		f _i	%	f _i	%	f _i	%	f _i	%	f _i	%		
1	152	67,65	164	51,25	219	55,58	134	46,53	59	60,20	77	45,83	
2	65	28,89	105	32,81	125	31,73	103	35,76	27	27,55	51	30,36	
3	6	2,67	35	10,94	32	8,12	36	12,50	10	10,20	26	15,48	
4	2	0,89	15	4,69	15	3,81	13	4,51	2	2,04	10	5,95	
5	-	-	1	0,31	3	0,76	2	0,69	-	-	4	2,38	
6	-	-	-	-	-	-	-	-	-	-	-	-	
7	-	-	-	-	-	-	-	-	-	-	-	-	
8	-	-	-	-	-	-	-	-	-	-	-	-	
n	225		320		394		288		98		168		
$\bar{x}$	1,3689		1,7000		1,6320		1,7708		1,5408		1,8869		
s ²	0,3395		0,7475		0,6920		0,7808		0,5749		1,0527		
V	0,4256		0,5086		0,5097		0,4990		0,4921		0,5438		

Fortsetzung Tabelle 12.28

Silben	B r i e f N r.									
	641		605		644		659		647	
	f _i	%	f _i	%	f _i	%	f _i	%	f _i	%
1	266	58,46	428	59,86	338	47,86	151	62,14	259	57,05
2	133	29,23	215	30,07	242	34,28	68	27,98	132	29,06
3	44	9,67	51	7,13	76	10,77	16	6,58	37	8,15
4	11	2,42	19	2,66	36	5,10	7	2,88	19	4,19
5	1	0,22	1	0,14	11	1,56	1	0,41	6	1,32
6	-	-	1	0,14	1	0,14	-	-	1	0,22
7	-	-	-	-	-	-	-	-	-	-
8	-	-	-	-	2	0,28	-	-	-	-
n	455		715		706		243		454	
$\bar{x}$	1,5670		1,5357		1,8003		1,5144		1,6432	
s ²	0,6103		0,5956		1,0154		0,6037		0,8462	
V	0,4985		0,5025		0,5597		0,5131		0,5598	

Wie deutlich der Unterschied zwischen dem "Erlkönig" und zumindest einigen Briefen ausgeprägt ist, zeigt ein Vergleich zwischen "Erlkönig" und Brief 596 mit Hilfe des χ^2 -Tests. Wir erhalten $\chi^2_3 = 51,17$ und $\phi = 0,35$. Wie die beiden Werte zeigen, besteht ein erheblicher Unterschied zwischen den beiden Texten. Die Zahl der Silben pro Wort ist somit bei GOETHE eine objektive Stilcharakteristik (vgl. Kap. 11.2.2).

Die Tabelle 12.29 zeigt die Anzahl der Silben pro Wort in zwei lateinischen Verstexten, in zwei lateinischen Prosatexten und in zwei deutschen Balladen. Wie ich bereits in Kapitel 6.2 angedeutet habe, liegt in metrischen Texten die durchschnittliche Wortlänge häufig erheblich unter dem in der jeweiligen Sprache üblichen Wert.¹⁷ Das gleiche gilt auch für die Varianz der Anzahl der Silben pro Wort. Ein deutlicher Beleg hierfür sind die in der Tabelle 12.29 aufgeführten Verstexte. Sowohl der Mittelwert als auch die Varianz ist in den beiden Balladentexten deutlich niedriger als in der Mehrzahl der in Tabelle 12.28 aufgeführten Briefe. Das gleiche gilt auch für die beiden lateinischen Verstexte im Verhältnis zu den beiden lateinischen Prosatexten. Daß der Unterschied zwischen Verstexten und Prosatexten nicht nur zufälliger Natur ist, zeigt z.B. ein Vergleich der Mittelwerte von SALLUST und HORAZ. Der t-Test ergibt einen Wert von 5,2981. Der Unterschied zwischen den beiden Mittelwerten ist hochsignifikant.

Der Grund für die spezifische Verteilung der Anzahl der Silben pro Wort in Verstexten dürfte fast stets im Einfluß der Variablen 'Metrum' zu sehen sein. Im "Erlkönig" z.B. beträgt die mittlere Silbenzahl pro Zeile 9,625. Jeweils vier Silben einer Zeile müssen markiert sein. Unter diesen Bedingungen lassen sich nur mit Mühe längere Wörter verwenden. Auch in einem Gedicht wie "Der Totentanz", in dem zwischen zwei Hebungen maximal zwei Senkungen zu finden sind und eine Zeile maximal vier Hebungen aufweist, dürfte bereits die Verwendung von viersilbigen Wörtern mit erheblichen Schwierigkeiten verbunden sein. Es ist deswegen wenig erstaunlich, daß dort keine Wörter mit mehr als vier Silben zu finden sind.

Tab. 12.29: Verteilung der Anzahl der Silben pro Wort in zwei lateinischen Verstexten, zwei lateinischen Prosatexten und zwei deutschen Balladen

Silben	LUKREZ	HORAZ	CAESAR	SALLUST	GOETHE	SCHILLER
	f_i	f_i	f_i	f_i	f_i	f_i
	(1) %	(2) %	(3) %	(4) %	(5) %	(6) %
1	1575 21,91	705 22,82	184 23,74	122 17,38	218 63,74	580 58,12
2	2606 36,24	1075 34,80	204 26,32	249 35,47	99 28,95	296 29,66
3	2252 31,32	980 31,73	194 25,03	196 27,92	21 6,14	97 9,72
4	636 8,85	291 9,42	125 16,13	110 15,67	4 1,17	24 2,40
5	111 1,54	38 1,23	54 6,97	24 3,42	-	1 0,10
6	10 0,14	-	13 1,68	1 0,14	-	-
7	-	-	1 0,13	-	-	-
n	7190	3098	775	702	342	998
$\bar{x}$	2,3231	2,3143	2,6181	2,5271	1,4474	1,5671
s^2	0,9442	0,9342	1,6786	1,1325	0,4402	0,5962
V	0,4183	0,4176	0,4949	0,4211	0,4584	0,4927

(1) "De Rerum Natura" (nach der Zählung von OTT 1974)

(2) "De Arte Poetica" (nach der Zählung von OTT 1970)

(3) "De Bello Gallico"

(4) "Bellum Iugurthinum"

(5) "Der Totentanz" (Gesamttext)

(6) "Die Kraniche des Ibycus" (Gesamttext)

Es besteht jedoch nicht nur ein Unterschied zwischen metrischen Texten und Prosatexten hinsichtlich der Anzahl der Silben pro Wort, sondern auch Prosatexte können, wie schon GOETHEs "Briefe" zeigen (vgl. Tab. 12.28), erhebliche Unterschiede in der Verteilung der Anzahl der Silben pro Wort aufweisen. Die Hauptursache dürfte in der Variablen 'Wortfrequenz' zu sehen sein. Denn wie bereits G.K. ZIPF für eine Reihe von Sprachen gezeigt hat, besteht ein Zusammenhang zwischen Wortfrequenz und Wortlänge: Je häufiger ein Wort ist, desto kürzer ist es in der Regel auch. Der Zusammenhang zwischen Wortfrequenz und Wortlänge läßt sich an einer Reihe sprachlicher Erscheinungen belegen. So besteht z.B. bei einer soziokulturell bedingten Erhöhung der Frequenz eines langen Wortes die Tendenz, dieses Wort zu kürzen oder durch ein kürzeres Wort zu ersetzen (vgl. ZIPF 1965:21ff).

Auf der Textebene hängt die Wortfrequenz wiederum von Variablen wie Kommunikationsintention, Thematik oder intendierter Leser ab. Das kann dazu führen, daß in der Sprache seltene und damit lange Wörter in einem Text überdurchschnittlich häufig verwendet werden. Dies dürfte auch der Grund dafür sein, daß in Sachprosa die mittlere Wortlänge meist erheblich höher ist als in 'normaler' Prosa (vgl. z.B. die Angaben bei FUCKS/LAUTER 1965:114f). Ein weiterer Beleg für den genannten Zusammenhang sind GOETHEs "Briefe". Briefe, die an gute Freunde gerichtet sind und die alltägliche Dinge zum Thema haben, scheinen insgesamt gesehen eine geringere mittlere Wortlänge aufzuweisen als z.B. 'Geschäftsbriefe' an Adressaten, zu denen ein sozial-distanziertes Verhältnis besteht.¹⁸

Versucht man die Verteilung der Anzahl der Silben pro Wort durch eine theoretische Verteilung zu beschreiben, wird ebenfalls der Unterschied zwischen Vers- und Prosatexten deutlich. Während für viele Prosatexte verschiedener Sprachen das Modell einer verschobenen negativen Binomialverteilung eine sehr gute Anpassung ergibt, scheint für Verstexte in vielen Fällen eher eine verschobene (positive) Binomialverteilung oder eine verschobene Poisson-Verteilung in Frage zu kommen (vgl. GROTHJAHN 1977).

12.7 Analyse der Lautstruktur

12.7.1. In den Kapiteln 5 und 10 habe ich darauf hingewiesen, daß aufgrund einer spezifischen Distribution der Laute eines Textes lautliche Wirkungen entstehen können und daß auch die rhythmische Struktur eines Textes durch die Lautdistribution beeinflusst werden kann.

Ein relativ häufig verwendetes Modell zur quantitativen Charakterisierung der Lautstruktur von Texten ist das informationstheoretische Modell der Entropie. Bei diesem Modell geht man davon aus, daß zwischen der Auftretenswahrscheinlichkeit eines beliebigen Zeichens - in unserem Fall der Laute eines Textes - und dem jeweiligen Informationsgehalt des Zeichens die Beziehung

$$I = \lg \frac{1}{p_i} = - \lg p_i$$

besteht, wobei $\lg$ für $\log_2$ steht, d.h. den dyadischen (binären) Logarithmus bezeichnet. Betrachtet man nun I als eine Zufallsvariable, dann beträgt ihr Erwartungswert

$$E(I) = H_1 = - \sum_i p_i \lg p_i \quad (12.97)$$

(12.97) ist ein Maß für den mittleren Informationsgehalt eines endlichen Zeicheninventars und wird als Entropie (1. Ordnung) oder wegen des negativen Vorzeichens auch als Negentropie bezeichnet.

Will man nicht mit dyadischen, sondern z.B. mit dekadischen oder natürlichen Logarithmen rechnen, kann man folgende Umrechnungsformeln benutzen:

$$\lg x = \frac{\log_{10} x}{\log_{10} 2} \quad (12.98a)$$

$$\lg x = \frac{\log_e x}{\log_e 2} \quad (12.98b)$$

Diese beiden Formeln sind lediglich ein Spezialfall der allgemeinen Umrechnungsformel für Logarithmen beliebiger Basis

$$\log_b x = \log_a x \cdot \log_b a, \quad (12.98)$$

die man leicht durch folgende Überlegung erhält:

Laut Definition des Logarithmus ist $a^y = x \Leftrightarrow \log_a x = y$.

Folglich gilt

$$x = a^{\log_a x}.$$

Logarithmiert man nun zur Basis b, folgt unmittelbar (12.98).

Die von SHANNON/WEAVER (1949) in die Informationstheorie eingeführte Entropie ist als Stilcharakteristik zum ersten Mal von FUCKS (1953) verwendet worden (vgl. FUCKS 1968:90). Die Entropie ist wie die Wiederholungsrate (vgl. Kap. 12.4.2) ein Maß für die Gleichverteilung von Häufigkeiten. Sind alle Häufigkeiten gleich, erreicht H_1 sein Maximum. Bezeichnet man den Umfang des jeweiligen Zeicheninventars mit k, dann gilt bei einer Gleichverteilung

$$p_1 = p_2 = \dots = p_k = \frac{1}{k},$$

und (12.97) wird zu

$$H_{1,\max} = H_0 = - \sum_{i=1}^k \frac{1}{k} \lg \frac{1}{k} = -k \frac{1}{k} (\lg 1 - \lg k) = \lg k. \quad (12.99)$$

H_1 erreicht sein Minimum, wenn die Wahrscheinlichkeit eines Zeichens 1 und die der übrigen Zeichen 0 beträgt. In diesem Fall gilt

$$H_{1,\min} = -1 \lg 1 = 0.$$

H_1 bewegt sich somit im Intervall $(0, \lg k)$. Wie man sieht, hängt das Intervall vom Inventarumfang ab. Dividiert man H_1 durch $H_{0,\max}$, erhält man die relative Entropie h mit dem Wertebereich $(0,1)$:

$$h = \frac{- \sum_i p_i \lg p_i}{\lg k} \quad (12.100)$$

Das Komplement zu h, d.h. $1 - h$ wird als relative Redundanz bezeichnet:

$$Re = 1 - h \quad (12.101)$$

Es soll nun die Entropie der Phoneme des "Erk König" berechnet werden. Dazu habe ich den Text des "Erk König" transkribiert und die Häufigkeiten der insgesamt 873 Phoneme ermittelt.¹⁹ Die Häufigkeiten sind in der Tabelle 12.30 aufgeführt. Zum Vergleich sind die Ergebnisse der umfangreichen Zählung von MEIER (1967:252f) angegeben, der 48.222 Laute aus Prosatexten und 48.143 Laute aus poetischen Texten ausgezählt hat.

Die durchgeführte Transkription, die weitgehend phonologisch ist, soll nur kurz diskutiert werden. Die Ermittlung des Phoneminventars beruht im wesentlichen auf UNGEHEUER (1969). Gegen UNGEHEUER nehme ich ein Phonem /ə/ an, anstatt [ə] als ein Allophon von /e/ zu werten. Die Affrikaten werden mit MORCINIEC (1958) biphonematisch, die Diphthonge abweichend monophonematisch gewertet. Außerdem werden /x/ und /ç/ als eigenständige Phoneme angesehen. Dieses Vorgehen läßt sich auch phonologisch rechtfertigen, da /x/ und /ç/ bei Berücksichtigung der Morphemstruktur deutscher Wörter zwar in komplementärer Distribution stehen, bei Nicht-Berücksichtigung der Morphemstruktur jedoch Oppositionen wie /ku:xən/ ~ /ku:çən/ existieren. Der Hauptgrund für die eigenständige Wertung von /ə, ç, x/ ist aber darin zu sehen, daß /ə/ und /e/ und auch /x/ und /ç/ phonetisch nicht sehr 'ähnlich' sind - auf die Problematik der phonetischen 'Ähnlichkeit' soll hier nicht eingegangen werden - und somit deutlich die Klanggestalt des Textes mitprägen. Die eigenständige Klassifikation von /ə, ç, x/ und die monophonematische Wertung der Diphthonge hat darüber hinaus noch den Vorteil, daß die Häufigkeiten von /ə, ç, x/ und /ae, ao, ɔɪ/

Tab. 12.30: Häufigkeiten der Phoneme des Deutschen nach MEIER (1967:251ff) und Häufigkeiten der Phoneme des "Erlkönig"

Phonem	MEIER (1967)				Erlkönig	
	Prosa		Poesie			
	n _i	%	n _i	%	n _i	%
r	3750	7,777	3534	7,341	66	7,560
l	1747	3,623	2147	4,460	27	3,093
m	1287	2,669	1577	3,276	41	4,696
n	5092	10,559	5183	10,766	111	12,715
ŋ	347	0,720	316	0,656	2	0,229
h	540	1,120	586	1,217	7	0,802
f	1182	2,451	1109	2,304	17	1,947
v	969	2,009	1238	2,572	17	1,947
s	2314	4,799	2158	4,482	27	3,093
z	1133	2,350	1012	2,102	23	2,635
ʃ	681	1,412	853	1,772	13	1,489
ç	1114	2,310	1050	2,181	34	3,895
x	332	0,689	389	0,808	2	0,229
j	105	0,218	92	0,191	1	0,115
b	921	1,910	830	1,724	14	1,604
p	379	0,786	386	0,802	4	0,458
d	2355	4,884	2121	4,406	33	3,780
t	4085	8,471	4538	9,426	97	11,111
g	1100	2,281	1014	2,106	14	1,604
k	796	1,651	928	1,928	15	1,718
ʒ	-	-	1	0,002	-	-
a:	899	1,864	707	1,469	16	1,833
a	1445	2,997	1411	2,931	25	2,864
e:	1248	2,588	1116	2,318	21	2,405
e	1306	2,708	1158	2,405	17	1,947
ɛ:	134	0,278	108	0,224	1	0,115
ə	4818	9,991	3848	7,993	66	7,560

Fortsetzung Tab. 12.30

Phonem	n_i	%	n_i	%	n_i	%
i:	1377	2,855	1322	2,746	15	1,718
ɪ	1842	3,820	2120	4,404	51	5,842
o:	506	1,049	533	1,107	13	1,489
ɔ	595	1,234	693	1,439	4	0,458
ø:	118	0,245	108	0,224	9	1,031
œ	58	0,120	44	0,091	3	0,344
u:	522	1,083	550	1,142	10	1,145
ʊ	996	2,065	1113	2,312	12	1,375
y:	198	0,410	249	0,517	3	0,344
ʏ	153	0,317	144	0,299	2	0,229
æ	1234	2,559	1224	2,542	35	4,009
ao	384	0,796	468	0,972	4	0,458
ɔ ø	160	0,332	165	0,343	1	0,115
	48222		48143		873	

auch nachträglich noch zusammengefaßt bzw. den entsprechenden Vokalen zugerechnet werden können. Das gewählte Vorgehen führt dadurch auch zu einer besseren Vergleichsmöglichkeit mit anderen Phonemzählungen (vgl. KING 1966). MEIER (1967) unterscheidet bei den Vokalen noch zwischen geschlossenen langen Vokalen und geschlossenen kurzen Vokalen und bei den Auslautkonsonanten zwischen stimmlosen Konsonanten und desonorisierten stimmhaften Konsonanten (z.B. /d/). In der Tabelle 12.30 sind diese Klassen nicht einzeln aufgeführt, sondern als geschlossene lange Vokale und stimmhafte Konsonanten zusammengefaßt.

Bei der Berechnung der Entropie H_1 werden die Wahrscheinlichkeiten p_i aus Formel (12.97) durch die relativen Häufigkeiten $\hat{p}_i = n_i/n$ abgeschätzt. H_1 kann somit auch direkt aus den absoluten Häufigkeiten berechnet werden. In diesem Fall wird (12.97) zu

$$H_1 = - \sum_i \frac{n_i}{n} \lg \frac{n_i}{n} = \lg n - \frac{1}{n} \sum_i n_i \lg n_i \quad (12.102)$$

Für den "Erlkönig" erhalten wir

$$\begin{aligned} \hat{H}_1 &= \lg 873 - \frac{1}{873} (66 \lg 66 + 27 \lg 27 + \dots + 1 \lg 1) \\ &= 4,5628. \end{aligned}$$

Da im "Erlkönig" 39 verschiedene Phoneme vorkommen, ergibt sich nach (12.100) für die relative Entropie

$$\hat{h} = \frac{4,5628}{\lg 39} = 0,8633.$$

Entsprechend erhalten wir für das Korpus von MEIER (1967):

$$\text{Prosa:} \quad \hat{H}_1 = 4,6909 \quad \hat{h} = \frac{4,6909}{\lg 39} = 0,8875$$

$$\text{Poesie:} \quad \hat{H}_1 = 4,7155 \quad \hat{h} = \frac{4,7155}{\lg 40} = 0,8861$$

Die Werte der relativen Entropien für das Prosa- und das Poesie-Korpus von MEIER weisen so gut wie keinen Unterschied auf. Etwas niedriger ist jedoch der Wert für den "Erlkönig". Ein deutlicherer Unterschied zwischen Prosa und Vers zeigt sich dagegen im Tschechischen, für das DOLŽEL (1963) folgende Entropien erster Ordnung ermittelt hat:

Vers: 4,5722
Prosa: 4,5919 bis 4,7044

LEVÝ (1965:216) interpretiert dieses Ergebnis folgendermaßen: "Die kleine Entropie der 1. Ordnung spricht dafür, daß seltene Laute (Buchstaben) in der Poesie noch seltener sind als in der Prosa und daß das Übergewicht der häufigen Laute im Vers noch markanter ist. Der Vers verwertet also die für die betreffende Sprache typischen Laute viel mehr und unterdrückt die seltenen."

Ein dem Tschechischen analoges Ergebnis scheint FÓNAGY (1961) für das Ungarische gefunden zu haben. FÓNAGY hat Texte laut für laut erraten lassen und hat dabei festgestellt, daß in der Poesie die Prädiktabilität der Phoneme wesentlich größer ist als in der Prosa, daß also die Poesie auf der Phonemebene redundanter als die Prosa ist. Auf der Lexemebene ist dagegen die Redundanz im Gedicht geringer als z.B. in einem Leitartikel, da im Gedicht z.B. attributive oder adverbiale Konstruktionen oder auch Vergleiche weit weniger gebunden sind als in Texten mit primär kommunikativer Funktion. Es ist m.E. allerdings fraglich, ob das von FÓNAGY auf diese Weise erzielte Ergebnis mit der mathematisch definierten Entropie direkt vergleichbar ist.

12.6.2. Es soll nun untersucht werden, ob die von LEVÝ (1965) gegebene Interpretation der Entropie 1. Ordnung auch auf den "Erlkönig" zutrifft. Hierzu testen wir zuerst den Unterschied zwischen der Entropie des "Erlkönig" ($\hat{H}_{1E}$) und der des Prosa-Korpus von MEIER ($\hat{H}_{1M}$), indem wir eine Transformation auf die Normalverteilung durchführen. Unter der Annahme der Unabhängigkeit von $\hat{H}_{1E}$ und $\hat{H}_{1M}$ lautet das Testkriterium

$$\hat{z} = \frac{\hat{H}_{1E} - \hat{H}_{1M} - [E(\hat{H}_{1E}) - E(\hat{H}_{1M})]}{\sqrt{V(\hat{H}_{1E}) + V(\hat{H}_{1M})}} \quad (12.103)$$

Den Erwartungswert und die Varianz der Zufallsvariablen $\hat{H}_1$ erhalten wir folgendermaßen (vgl. BASHARIN 1959, ALTMANN/LEHFELDT 1979):

Verwendet man anstelle von dyadischen Logarithmen natürliche Logarithmen, lautet $\hat{H}_1$ (vgl. 12.98)

$$\hat{H}_1 = H_1(\hat{p}_i) = - \sum_i \hat{p}_i \ln \hat{p}_i \ln e.$$

Wir entwickeln nun

$$\frac{H_1(\hat{p}_i)}{\ln e} = - \sum_i \hat{p}_i \ln \hat{p}_i$$

an der Stelle $E(\hat{p}_i) = p_i$ in eine Taylorreihe (vgl. KENDALL/STUART I 1969:231ff). Da die Ableitungen der Ordnung $r \geq 3$ zu Ausdrücken der Größenordnung kleiner oder gleich n^{-2} führen (vgl. BASHARIN 1959), beschränken wir uns auf die 1. und 2. Ableitung:

$$\begin{aligned} \frac{H_1(\hat{p}_i)}{\ln e} &= H_1(p_i) + \left[\sum_i (\hat{p}_i - p_i) \frac{\partial \frac{H_1(\hat{p}_i)}{\ln e}}{\partial \hat{p}_i} \right]_{p_i} \\ &+ \frac{1}{2!} \left[\sum_i (\hat{p}_i - p_i)^2 \frac{\partial^2 \frac{H_1(\hat{p}_i)}{\ln e}}{\partial \hat{p}_i^2} \right]_{p_i} \\ &+ \sum_{i \neq j} \sum (\hat{p}_i - p_i)(\hat{p}_j - p_j) \frac{\partial^2 \frac{H_1(\hat{p}_i)}{\ln e}}{\partial p_i \partial p_j} \bigg|_{p_i p_j} \end{aligned}$$

Da

$$\frac{\partial \frac{H_1(\hat{p}_i)}{\ln e}}{\partial \hat{p}_i} \bigg|_{p_i} = \frac{\partial (- \sum \hat{p}_i \ln \hat{p}_i)}{\partial \hat{p}_i} \bigg|_{p_i} = - (1 + \ln p_i)$$

und

$$\frac{\partial^2 \frac{H_1(\hat{p}_i)}{\ln e}}{\partial \hat{p}_i^2} \bigg|_{p_i} = - \frac{1}{p_i}$$

ist und da außerdem die gemischten Ableitungen im Punkte $(p_1, \dots, p_k)$ verschwinden, lauten die ersten Glieder der Taylorreihe:

$$\frac{\hat{H}_1}{\ln e} = \frac{H_1}{\ln e} - \sum_{i=1}^k (\hat{p}_i - p_i)(1 + \ln p_i) - \frac{1}{2} \sum_{i=1}^k \frac{(\hat{p}_i - p_i)^2}{p_i} \quad (12.104)$$

Wir berechnen zuerst den Erwartungswert von (12.104):

$$E\left(\frac{\hat{H}_1}{\ln e}\right) = \frac{H_1}{\ln e} - \sum_i (1 + \ln p_i) E(\hat{p}_i - p_i) - \frac{1}{2} \sum_i \frac{E(\hat{p}_i - p_i)^2}{p_i}$$

Berücksichtigt man, daß

$$E(\hat{p}_i - p_i) = 0$$

ist und für multinomialverteilte Häufigkeiten

$$V(\hat{p}_i) = E(\hat{p}_i - p_i)^2 = \frac{p_i(1 - p_i)}{n}$$

gilt, erhält man

$$E\left(\frac{\hat{H}_1}{\ln e}\right) = \frac{H_1}{\ln e} - \frac{1}{2n} \sum_{i=1}^k \frac{p_i(1 - p_i)}{p_i} = \frac{H_1}{\ln e} - \frac{1}{2n} (k - 1).$$

Multipliziert man nun mit $\ln e$, bekommt man den gesuchten Erwartungswert:

$$E(\hat{H}_1) = H_1 - \frac{k-1}{2n} \ln e + O\left(\frac{1}{n^2}\right) \quad (12.105)$$

Bei der Berechnung der Varianz von $\hat{H}_1$ gehen wir wiederum von der Taylorreihe (12.104) aus. Wir berücksichtigen diesmal jedoch nur die beiden ersten Glieder, da die Einbeziehung weiterer Glieder lediglich zu Ausdrücken kleiner oder gleich $O(n^{-2})$ führt (vgl. BASHARIN 1959). Subtrahiert man auf beiden Seiten $H_1/\ln e$, quadriert dann und bildet schließlich den Erwartungswert, erhält man

$$E\left(\frac{\hat{H}_1}{\ln e} - \frac{H_1}{\ln e}\right)^2 = E\left[\sum_i (\hat{p}_i - p_i)(1 + \ln p_i)\right]^2.$$

Da

$$E\left(\frac{\hat{H}_1}{\ln e} - \frac{H_1}{\ln e}\right)^2 = \frac{E(\hat{H}_1 - H_1)^2}{\ln^2 e} = \frac{V(\hat{H}_1)}{\ln^2 e}$$

ist, folgt

$$\begin{aligned} \frac{V(\hat{H}_1)}{\ln^2 e} &= E\left[\sum_i (\hat{p}_i - p_i)(1 + \ln p_i)\right]^2 \\ &= E\left[\sum_i (\hat{p}_i - p_i)^2(1 + \ln p_i)^2\right. \\ &\quad \left.+ \sum_{i \neq j} (\hat{p}_i - p_i)(\hat{p}_j - p_j)(1 + \ln p_i)(1 + \ln p_j)\right]. \end{aligned}$$

Berücksichtigt man, daß

$$E(\hat{p}_i - p_i)^2 = V(\hat{p}_i) = \frac{p_i(1-p_i)}{n}$$

und

$$E(\hat{p}_i - p_i)(\hat{p}_j - p_j) = \text{Cov}(\hat{p}_i, \hat{p}_j) = -\frac{p_i p_j}{n}$$

ist, erhält man

$$\frac{V(\hat{H}_1)}{\ln^2 e} = \sum_i (1 + \ln p_i)^2 \frac{p_i(1-p_i)}{n} - \sum_{i \neq j} (1 + \ln p_i)(1 + \ln p_j) \frac{p_i p_j}{n}. \quad (12.106)$$

Um die Doppelsumme besser auswerten zu können, erweitern wir den Summationsbereich auch auf den Fall $i=j$. Die dadurch erhaltene zusätzliche $\sum_i (1 + \ln p_i)^2 p_i^2$ ziehen wir wieder ab:

$$\begin{aligned} \frac{V(\hat{H}_1)}{\ln^2 e} &= \frac{1}{n} \sum_i (1 + \ln p_i)^2 p_i(1 - p_i) \\ &\quad - \frac{1}{n} \sum_{i \neq j} (1 + \ln p_i)(1 + \ln p_j) p_i p_j + \frac{1}{n} \sum_i (1 + \ln p_i)^2 p_i^2 \\ &= \frac{1}{n} \left[\sum_i (1 + \ln p_i)^2 p_i - \sum_i (1 + \ln p_i)^2 p_i^2 \right. \\ &\quad \left. - \left(\sum_i (1 + \ln p_i) p_i \right)^2 + \sum_i (1 + \ln p_i)^2 p_i^2 \right] \\ &= \frac{1}{n} \left[\sum_i p_i \ln^2 p_i - \left(\sum_i p_i \ln p_i \right)^2 \right] \quad (12.107) \end{aligned}$$

Multipliziert man nun beide Seiten von (12.107) mit $\ln^2 e$, erhält man schließlich das gesuchte Ergebnis:

$$V(\hat{H}_1) = \frac{1}{n} \left[\sum_i p_i \ln^2 p_i - H_1^2 \right] + O\left(\frac{1}{n^2}\right) \quad (12.108)$$

12.6.3. Mit Hilfe der abgeleiteten Formeln können wir nun testen, ob sich "Erlkönig" und Prosa-Korpus von MEIER hinsichtlich der Entropie 1. Ordnung unterscheiden. Dazu schätzen wir wiederum die einzelnen p_i (und damit auch H_1) direkt aus den Daten. Eine andere Möglichkeit bestände darin, das von ALTMANN/LEHFELDT (1979) vorgeschlagene Modell einer geometrischen Verteilung zur Schätzung der p_i zu benutzen.

Da unter der Nullhypothese $E(\hat{H}_{1E}) = E(\hat{H}_{1M})$ gilt, können wir auf die Berechnung der Erwartungswerte verzichten. Für die Varianzen erhalten wir unter Benutzung von (12.108) folgende Werte:

Prosa-Korpus von MEIER:

$$V(\hat{H}_{1M}) = \frac{1}{48222} (23,48 - 4,69^2) = 0,000029$$

"Erlkönig":

$$V(\hat{H}_{1E}) = \frac{1}{873} (22,43 - 4,56^2) = 0,00184$$

Der Normalverteilungstest ergibt

$$\hat{z} = \frac{|4,5628 - 4,6909|}{\sqrt{0,00184 + 0,000029}} = 2,9608.$$

Da $P(|z| \geq 2,9608) \leq 0,003$ ist, kann mit einer Irrtumswahrscheinlichkeit von 0,3 % angenommen werden, daß zwischen den beiden Entropien ein Unterschied besteht. Da jedoch der durchgeführte Test lediglich approximativen Charakter hat, gilt dieses Testresultat nur mit Einschränkungen. Vergleichen wir auch noch die beiden Korpora von MEIER untereinander, so erhalten wir für die Varianz des Poesie-Korpus $V(\hat{H}_1) = 0,0000277$. Der z-Test ergibt einen Wert von 0,00755. Es besteht somit kein Unterschied zwischen den Entropien des Prosa-Korpus und des Poesie-Korpus.

Wir wollen nun die Verteilung der Phonemhäufigkeiten im "Erlkönig" und im Prosa-Korpus noch zusätzlich mit Hilfe eines χ^2 -Tests auf Homogenität prüfen. Wir benutzen wiederum die Formel von BRANDT/SNEDECOR (vgl. Formel 12.29):

$$\chi^2 = \frac{49095^2}{48222(873)} \left(\frac{66^2}{3816} + \frac{27^2}{1774} + \dots + \frac{1^2}{161} - \frac{873^2}{49095} \right) = 114,87$$

Um die beiden Testergebnisse vergleichen zu können, benötigen wir für den erhaltenen χ^2 -Wert die exakte Überschreitungswahrscheinlichkeit. Die meisten χ^2 -Tafeln geben jedoch lediglich kritische Schranken an. Wir benutzen deshalb zur Berechnung der Überschreitungswahrscheinlichkeit eine Approximation durch die Normalverteilung. Da die Transformation

$$z = \frac{\chi^2 - E(\chi^2)}{\sqrt{V(\chi^2)}} = \frac{\chi^2 - v}{\sqrt{2v}}$$

erst bei einer sehr großen Anzahl von Freiheitsgraden hinlänglich genaue Werte liefert und auch die übliche auf R.A. FISHER zurückgehende Transformation

$$z = \sqrt{2\chi^2} - \sqrt{2v - 1}$$

bei $FG = v = 38$ noch relativ ungenau ist, verwenden wir die WILSON-HILFERTY-Approximation (vgl. JOHNSON/KOTZ 1970:176). Wie WILSON und HILFERTY gezeigt haben, folgt $(\chi^2/v)^{1/3}$ asymptotisch einer Normalverteilung $N(1 - 2/9v, 2/9v)$. Es gilt somit

$$z = \frac{\left(\frac{\chi^2}{v}\right)^{\frac{1}{3}} + \frac{2}{9v} - 1}{\sqrt{\frac{2}{9v}}} \quad (12.109)$$

Wir erhalten

$$\hat{z} = \frac{\left(\frac{114,87}{38}\right)^{\frac{1}{3}} + \frac{2}{9(38)} - 1}{\sqrt{\frac{2}{9(38)}}} = 5,9075.$$

Da diesem Wert eine Überschreitungswahrscheinlichkeit von etwa $1,8 \cdot 10^{-9}$ entspricht, ist die Hypothese der Homogenität der beiden Häufigkeitsverteilungen abzulehnen.

Im Vergleich zum Entropie-Test weist der χ^2 -Test auf einen weit stärker ausgeprägten Unterschied zwischen den beiden Häufigkeitsverteilungen hin. Es ist jedoch zu berücksichtigen, daß beide Tests unterschiedliche Eigenschaften testen. Würden z.B. die beiden Häufigkeitsverteilungen identische Häufigkeiten aufweisen, allerdings in genau umgekehrter Reihenfolge, dann würde der Unterschied in der Rangfolge durch den χ^2 -Test erfaßt, die Entropien beider Verteilungen wären hingegen trotz unterschiedlicher Rangfolge gleich.

12.6.4. Mit Hilfe des Entropie-Test ist zwar nachgewiesen, daß der "Erlkönig" und das Prosa-Korpus von MEIER in unterschiedlicher Weise von einer Gleichverteilung abweichen, es ist jedoch noch nicht nachgewiesen, daß im "Erlkönig" seltene Laute noch seltener und häufige Laute noch häufiger als in der Prosa sind. Um dies zu überprüfen, testen wir den Unterschied zwischen den Phonemhäufigkeiten des "Erlkönig" und des MEIER-Korpus mit Hilfe der Normalverteilung. Vernachlässigt man den relativ geringen Stichprobenfehler des MEIER-Korpus und betrachtet das Poesie- und das Prosa-Korpus jeweils als Grundgesamtheit,²⁰ dann lautet das Testkriterium

$$\hat{z} = \frac{\hat{p} - E(\hat{p})}{\sqrt{V(\hat{p})}} = \frac{\hat{p} - p}{\sqrt{\frac{p(1-p)}{n}}}$$

Beim Vergleich z.B. der Häufigkeit des Phonems /r/ im "Erlkönig" und im Poesie-Korpus erhalten wir (vgl. Tab. 12.30)

$$\hat{z} = \frac{7,56 - 7,341}{\sqrt{\frac{7,341(92,695)}{873}}} = 0,2481.$$

Was die festzusetzende Irrtumswahrscheinlichkeit betrifft, so stellt sich wie bereits in Kapitel 12.5.2 das Problem, daß die durchzuführenden Tests als eine Familie simultaner Tests aufgefaßt werden können. Um die Teststärke nicht zu sehr zu reduzieren, verzichten wir diesmal jedoch auf eine Adjustierung der Irrtumswahrscheinlichkeit und betrachten $\hat{z}$ -Werte größer als $|1,96|$ als signifikant. Die Ergebnisse des Vergleichs sind in Tabelle 12.31 aufgeführt. Signifikante Werte sind durch * gekennzeichnet.

Tab. 12.31: Vergleich der Phonemfrequenzen im "Erlkönig" und bei MEIER (1967:253)

Phonem	Prosa $\hat{z}$	Poesie $\hat{z}$	Phonem	Prosa $\hat{z}$	Poesie $\hat{z}$
r	-0,2394	0,2481	ʒ	-	-0,1321
l	-0,8380	-1,9567	a:	-0,0677	0,8939
m	3,7159*	2,3570*	a	-0,2305	-0,1174
n	2,0729*	1,8579	e:	-0,3405	0,1708
ŋ	-1,7159	-1,5628	ɛ	-1,3853	-0,8833
h	-0,8928	-1,1183	ɛ:	-0,9147	-0,6812
f	-0,9631	-0,7031	ø	-2,3952*	-0,4718
v	-0,1306	-1,1666	i:	-2,0172*	-1,8586
s	-2,3583*	-1,9835*	ɪ	3,1188*	2,0707*
z	0,5559	1,0978	o:	1,2760	1,0787
ʃ	0,1928	-0,6338	ɔ	-2,0769*	-2,3669*
ʒ	3,1175*	3,4672*	ø:	4,6976*	5,0436*
x	-1,6431	-1,9109	æ	1,9117	2,4792*
j	-0,6525	-0,5143	u:	0,1770	0,0083
b	-0,6605	-0,2724	ɔ	-1,4336	-1,8422
p	-1,0974	-1,1395	y:	-0,3052	-0,7127
d	-1,5134	-0,9012	ʏ	-0,4625	-0,3788
t	2,8013*	1,7039	ae	2,7131*	2,7539*
g	-1,3398	-1,0330	ao	-1,1238	-1,5480
k	0,1554	-0,4512	ɔø	-1,1146	-1,1522

Operationalisiert man die Begriffe 'seltener Laut' und 'häufiger Laut', indem man Laute (Phoneme), deren relative Häufigkeit (%) den unter der Hypothese einer Gleichverteilung zu erwartenden Wert von $100/39 = 2,56$ überschreitet, als 'häufig' und diejenigen Laute (Phoneme), deren Häufigkeit kleiner oder gleich diesem Wert ist, als 'selten' bezeichnet, dann kann die Tabelle 12.31 folgendermaßen interpretiert werden: Von den insgesamt 11 Phonemen, die beim Vergleich zwischen "Erlkönig" und Prosa-Korpus den kritischen Wert von 11,961 überschreiten, gehören 4 zu den seltenen und 7 zu den häufigen Phonemen. Unter den in Prosatexten seltenen Phonemen sind im "Erlkönig" 2 Phoneme (/æ, ɔ/) seltener und 2 Phoneme (/ç, ø:/) häufiger als in Prosa. Bei den in Prosa häufigen Phonemen sind im "Erlkönig" 4 Phoneme (/m, n, t, l/) häufiger und 3 Phoneme (/s, ə, i:/ weniger häufig als in Prosa. Es kann somit in bezug auf den "Erlkönig" wohl kaum von einer eindeutigen Tendenz zur Verringerung der Frequenz der seltenen Phoneme und Erhöhung der Frequenz der häufigen Phoneme gesprochen werden.²¹ Wie außerdem ein Vergleich der χ^2 -Werte für das Poesie- und das Prosa-Korpus zeigt, weisen beide Korpora fast stets bei den gleichen Phonemen Unterschiede zum "Erlkönig" auf. Auffallend in bezug auf den "Erlkönig" ist u.a. die hohe Frequenz der in Prosa und Poesie sehr seltenen Phoneme /æ/ und /ø:/. Die deutlich erhöhte Frequenz des /ø:/ kann dabei ebenso wie die erhöhte Frequenz des /ç/ auf das mehrmalige Vorkommen des Schlüsselworts 'Erlkönig' zurückgeführt werden.

12.6.5. Vergleicht man den "Erlkönig" und das Korpus von MEIER in bezug auf die Häufigkeit von Konsonanten und Vokalen, erhält man die Daten, die in Tabelle 12.32 aufgeführt sind. Wie die Tabelle zeigt, stimmen "Erlkönig" und Poesie-Korpus weitgehend überein. Es fällt auf, daß der Anteil der Konsonanten, der in deutschen Prosatexten im Vergleich zu Sprachen wie Italienisch oder auch Französisch bereits relativ groß ist, in der Poesie noch größer ist.²²

Tab. 12.32: Anzahl der Konsonanten und Vokale im "Erlkönig" und bei MEIER (1967:253)

	Erlkönig		MEIER (1967)			
	f_i	%	Prosa		Poesie	
	f_i	%	f_i	%	f_i	%
Konsonant	565	64,72	30230	62,69	31063	64,52
Vokal	308	35,28	17992	37,31	17080	35,48
	873		48222		48143	

Dies deckt sich auch mit den Ergebnissen von KOPCZYŃSKA/MAYENOWA (1970), die in polnischer Poesie ebenfalls eine im Verhältnis zur Prosa erhöhte Häufigkeit der Konsonanten festgestellt haben.

Der Unterschied zwischen Prosa und Poesie ist im Deutschen allerdings sehr gering. Vergleicht man den "Erlkönig" mit dem Prosa-Korpus, erhält man lediglich ein $\chi^2_1 = 1,51$ und damit ein nicht-signifikantes Ergebnis. Vergleicht man das Prosa- und Poesie-Korpus, ergibt sich zwar aufgrund des hohen Stichprobenumfangs ($n = 96.365$) ein hochsignifikanter Unterschied ($\chi^2_1 = 34,97$). Mißt man jedoch die Stärke der Assoziation zwischen den Untersuchungsvariablen mit Hilfe des Phi-Koeffizienten, erhält man lediglich $\phi = 0,019$. Berechnet man den Lambda-Koeffizienten (vgl. Kap. 12.2.5.2), ergibt sich sogar der Wert 0:

$$\lambda = \frac{E_1 - E_2}{E_1} = \frac{35072 - 35072}{35072} = 0$$

Dies bedeutet, daß die Kenntnis der Variablen 'Gattung' nicht zur Fehlerreduktion bei der Vorhersage des Verhältnisses von Vokal zu Konsonant im Deutschen beiträgt. Das gleiche gilt, wenn man die Variable 'Konsonant' bzw. 'Vokal' als unabhängige Variable ansieht.

Die Berechnung der Entropie und der Vergleich der einzelnen Phonemfrequenzen ist natürlich nur eine unter vielen Möglichkeiten zur quantitativen Charakterisierung der Lautstruktur von Texten. Weitere Möglichkeiten sind z.B.: Komponentenanalyse der Phoneme (vgl. SCHÄDLICH 1969; JONES 1969) und Berechnung der mittleren Phonemdistanz des Textes (vgl. ALTMANN 1969), Messung von Lautkorrespondenzen (vgl. z.B. SKINNER 1939; 1941; SEBEOK/ZEPS 1959; KNAUER 1965) oder Bestimmung der Euphonie (vgl. ALTMANN 1966a,b; HREBIČEK 1965).²³

Die bisher aufgezeigten Verfahren und Modelle, die nur eine Auswahl darstellen, dürften bereits deutlich gemacht haben, welche vielfältigen Möglichkeiten die Verwendung quantitativer Methoden eröffnet. Im folgenden Kapitel werde ich nun die jüngste Entwicklung in der Metrik, die generative Metrik, diskutieren. Dabei wird deutlich werden, daß auch diese Richtung erst durch eine Einbeziehung probabilistischer Modellvorstellungen zu einer adäquaten Metriktheorie weiterentwickelt werden kann.

13. DIE GENERATIVE METRIK

13.1. Der generative Ansatz in der Metrik

13.1.1. Die generative Metrik wurde Ende 1966 von M. HALLE und S.J. KEYSER mit dem Aufsatz "Chaucer and the study of prosody" begründet.¹ In dieser Arbeit wurde zum ersten Mal der Versuch unternommen, den Ansatz der generativen Linguistik, und zwar insbesondere von CHOMSKY (1965), systematisch auf die Metrik zu übertragen. Der Terminus 'generative Metrik' ist jedoch erst von BEAVER (1968a) verwendet worden.²

Konstitutiv für die Theorie von HALLE/KEYSER ist folgender Gedanke:³ Ein als Folge von Positionen verstandenes relativ invariantes, abstraktes metrisches Schema wird mit Hilfe von Zuordnungsregeln in unterschiedlicher Weise sprachlich realisiert. Die Zuordnungsregeln sind dabei so angeordnet, daß die Anwendung zugleich ein Maß für die Komplexität einer Verszeile abgibt.

Neben der Tatsache, daß in der Theorie von HALLE/KEYSER Begriffe wie Versfuß und Rhythmus keine Rolle spielen, unterscheidet sich diese Konzeption von vielen traditionellen Konzeptionen vor allem dadurch, daß Unterschiede in der sprachlichen Realisation eines Metrums - HALLE/KEYSER exemplifizieren ihre Theorie am jambischen Pentameter CHAUCERS - nicht auf unterschiedliche metrische Schemata, sondern auf Unterschiede in der Anwendung der Zuordnungsregeln zurückgeführt werden.

Der Ansatz von HALLE/KEYSER (1966) ist in einer Reihe von Veröffentlichungen einiger weniger Autoren weitergeführt worden.⁴ Obwohl dabei nicht unerhebliche Modifikationen an der ursprünglichen Theorie vorgenommen worden sind - auch HALLE/KEYSER selbst haben in ihrem Buch "English Stress: Its Form, its Growth, and its Role in Verse" (1971) die ursprüngliche Theorie in einigen Punkten revidiert - ist jedoch der Kern der Theorie von 1966 (relativ invariantes abstraktes metrisches Schema, relativ komplexe Zuordnungsregeln) nicht angetastet worden.⁵

Das Ziel der generativen Metrik wird von HALLE/KEYSER (1971) in der Aufstellung einer kohärenten und expliziten Theorie gesehen, die die Fähigkeit des Dichters und des erfahrenen ("experienced") Lesers erklärt, Verszeilen als metrisch oder unmetrisch und als mehr oder minder komplex zu beurteilen. Darüber hinaus soll eine solche Theorie die Beziehung dieser Fähigkeit zur normalen Sprachkompetenz erhellen und zum Verständnis des Wesens der metrischen Dichtung beitragen (vgl. HALLE/KEYSER 1971:139f). Die Affinität dieser Zielsetzung - eine ähnliche Zielvorstellung wird auch von anderen Vertretern der generativen Metrik genannt - zu den Zielen der generativen Transformationsgrammatik im Sinne von CHOMSKY (1965) ist evident.⁶

13.1.2. Die bisher skizzierte Entwicklung in der generativen Metrik ist vor allem durch den Versuch gekennzeichnet, den theoretischen Ansatz von HALLE/KEYSER (1966) zu exhaustieren. In jüngster Zeit sind jedoch von BOWLEY (1974) und von DEVINE/STEPHENS (1975) zwei weitere Versionen einer 'generativen' Metrik vorgeschlagen worden, die auch den Kern des bisherigen Ansatzes in Frage stellen.

Sowohl BOWLEY als auch DEVINE/STEPHENS gehen dabei - allerdings in unterschiedlicher Weise - von einer traditionellen Metrikkonzeption aus. Dies bedeutet u.a., daß BOWLEY und DEVINE/STEPHENS entgegen dem üblichen Ansatz in der generativen Metrik die (oberflächenstrukturelle) metrische Variation nicht auf Unterschiede in der Anwendung der Zuordnungsregeln (*correspondence rules*, *mapping rules*, *realization rules*), sondern auf unterschiedliche metrische Schemata zurückführen.

Der generative Ansatz BOWLEYS besteht nun lediglich darin, daß die Schemata durch analog zu den *phrase structure rules* der generativen Grammatik konzipierte metrische Strukturregeln generiert und mit Hilfe von nur zwei Korrespondenzregeln umkehrbar eindeutig sprachlich abgebildet werden.

Die Funktion dieser Regeln und die Hauptaufgabe der Metriktheorie werden von BOWLEY folgendermaßen charakterisiert:

"The main task of metrical theory therefore is to define the rules that govern underlying metrical patterns and to specify the correspondence rules that relate them to actual lines. The rules will thus determine what is metrical and what is not, just as grammatical rules determine what is grammatical or not." (1974:11)

DEVINE/STEPHENS (1975), die ihre Theorie nicht wie fast alle bisherigen Beiträge zur generativen Metrik an englischen, sondern primär an altgriechischen Beispielen exemplifizieren,⁷ halten den "item and arrangement approach" von BOWLEY für inakzeptabel und schlagen stattdessen einen transformationellen Ansatz vor, eine Möglichkeit, die BOWLEY im wesentlichen aus beschreibungstechnischen Gründen (Erhöhung der Komplexität des Regelapparates) ablehnt.

Die Konzeption von DEVINE/STEPHENS kann stark vereinfachend folgendermaßen beschrieben werden: Ein metrisches Basisschema - im Falle des klassischen Hexameters z.B. eine Folge von 6 Daktylen - wird mit Hilfe von metrischen Strukturregeln generiert und durch metrische Transformationsregeln, durch die Phänomene wie Zäsur, Katalexe, Anceps und Kontraktion berücksichtigt werden, in "metrical surface patterns" überführt. Die "metrical surface patterns" werden dann schließlich durch *mapping rules* sprachlich realisiert. Die metrischen Transformationsregeln leisten in dieser Konzeption somit einen Teil dessen, was bei BOWLEY die metrischen Strukturregeln und beim 'traditionellen' generativen Ansatz die Zuordnungsregeln bewirken.

DEVINE/STEPHENS begründen diese Konzeption, die zumindest in bezug auf den gewählten Untersuchungsgegenstand eine adäquatere Analyse ermöglicht als der übliche generative Ansatz, u.a. mit psychologischen Argumenten. So verweisen sie z.B. auf psychologische Untersuchungen, um gegenüber anderen Vertretern der generativen Metrik die Bedeutung des Versfußes als metrische Struktureinheit zu rechtfertigen. Dagegen spielen in den meisten Beiträgen zur generativen Metrik psychologische Argumente entweder überhaupt keine oder lediglich eine marginale Rolle, obwohl z.B. HALLE/KEYSER (1971) oder BEAVER (1974) den Anspruch erheben, die metrische Kompetenz des "experienced reader" bzw. des "perfect poet" zu erklären.

Geht man davon aus, daß eine Metriktheorie das tatsächliche Verhalten eines Autors bzw. Lesers abbilden soll, so scheint der Ansatz von DEVINE/STEPHENS einem solchen Anspruch zumindest eher gerecht zu werden als der traditionelle generative Ansatz.

13.1.3. Die Annahmen und Ziele der generativen Metrik sind von verschiedener Seite kritisiert worden. Sieht man einmal von der Detailkritik ab, dann scheinen mir zwei Kritikpunkte besonders wichtig zu sein.⁸

Der erste Kritikpunkt bezieht sich auf die erklärende bzw. prädiktive Kraft der generativen Metrik. So meinen z.B. BERNHART (1974), KLEIN (1974) und IHWE (1975), daß den Regeln von HALLE/KEYSER (1971) lediglich das Prädikat 'taxonomisch' zugeschrieben werden könne. Die gleiche Feststellung kann m.E. auch in bezug auf die '(erweiterte) ursprüngliche Theorie' (BEAVER 1974:7) und in bezug auf die Konzeptionen von BOWLEY und von DEVINE/STEPHENS getroffen werden.

Der zweite Kritikpunkt betrifft die Analogie zwischen 'metrischer Kompetenz' und Sprachkompetenz bei HALLE/KEYSER (1971) und anderen Vertretern der generativen Metrik. KLEIN (1974) kritisiert vor allem, daß es sich im Fall der metrischen Kompetenz nicht um eine implizite, sondern um eine explizite Fähigkeit handle. Diesem Einwand, der, wie auch BEAVER (1974) zugesteht, berechtigt ist, kann noch hinzugefügt werden, daß es zwar im Rahmen einer Grammatiktheorie zumindest für bestimmte Zielsetzungen durchaus vertretbar ist, von einer *s y m m e t r i s c h e n* Sprachkompetenz auszugehen, über die Sprecher und Hörer in gleicher Weise verfügen, daß jedoch in bezug auf die metrische Kompetenz eine analoge Annahme kaum gerechtfertigt sein dürfte (vgl. KINTGEN 1974).

Gegen eine analog zur Sprachkompetenz verstandene metrische Kompetenz kann jedoch noch ein weit schwerwiegenderer Einwand vorgebracht werden. Denn völlig inakzeptabel ist es m.E., wenn BEAVER (1974:26f) mit Bezug auf CHOMSKY (1965) feststellt:

"Just as a grammar supposes an ideal speaker in an ideal speech community with a never changing language, just so metrical rules presuppose the perfect poet in the ideal poetry-experiencing community with a never changing language."

Eine solche Idealisierung hat nämlich, sofern sie ernst genommen wird, zur Konsequenz, daß in der Metrik analog zur Grammatiktheorie CHOMSKYs von stilistisch-pragmatischen Faktoren zu abstrahieren wäre. Daß dies in der generativen Metrik bisher bereits weitgehend der Fall ist,⁹ wird deutlich von R. FOWLER (1976:26) festgestellt:

"The comparative neglect of context, register, motive and idiolect in Chomsky's grammar has a disturbing parallel in the neglect of genre, mode, purpose and individual style in generative metrics."

Eine adäquate Berücksichtigung stilistisch-pragmatischer Faktoren ist m.E. jedoch nur dann möglich, wenn die generative Metrik sich nicht länger lediglich am Paradigma CHOMSKYs orientiert, das aufgrund seiner deterministischen Konzeption für die Analyse nicht-idealisierten, tatsächlichen Sprachverhaltens nur bedingt geeignet ist, sondern auch solche neueren Ansätze innerhalb der generativen Grammatik berücksichtigt, die versuchen, mit Hilfe probabilistisch bewerteter Regeln zu einer adäquateren Analyse sprachlichen Verhaltens zu gelangen (vgl. hierzu KLEIN 1974a).

Es soll nun am Beispiel des lateinischen Hexameters gezeigt werden, wie die generative Metrik zu einer adäquateren Konzeption weiterentwickelt werden kann.

13.2. Generative Metrik und probabilistische Regelbewertung

13.2.1. CRUSIUS/RUBENBAUER (1963:48f) beschreiben den lateinischen daktylischen Hexameter folgendermaßen:

"Der *d a k t y l i s c h e* *H e x a m e t e r* besteht aus fünf vollständigen Daktylen und an sechster Stelle einem Spondeus oder Trochäus statt des Daktylus... In allen Füßen kann der Daktylus

durch den Spondeus ersetzt werden, am seltensten geschieht die Ersetzung im 5. Fuß."

In dieser Beschreibung - eine ähnliche Darstellung findet sich in vielen lateinischen Metriken - wird bereits darauf hingewiesen, daß der sog. *spondiacus*, d.h. ein Hexameter mit einem Spondeus im fünften Fuß, sehr selten ist. Aber auch in den ersten vier Füßen treten Daktylen und Spondeen nicht mit gleicher Häufigkeit auf, wie DROBISCH bereits 1866 anhand einer Stichprobe von 9493 Hexametern aus 15 lateinischen Dichtern gezeigt hat.

Ein solches Phänomen kann mit den Methoden der generativen Metrik bisher nicht in adäquater Weise erfaßt werden. Ein deutliches Beispiel hierfür ist die generative Darstellung des klassischen (griechischen) Hexameters bei HALLE (1970) und DEVINE/STEPHENS (1975).

Die Beschreibung HALLEs, die der traditionellen generativen Konzeption entspricht, kann - etwas vereinfacht - folgendermaßen wiedergegeben werden:

metrisches Schema: s w s w s w s w s s

Zuordnungsregel: Jedem w entspricht eine lange Silbe oder zwei kurze Silben. Jedem s entspricht eine lange Silbe. Die letzte Silbe eines Verses kann kurz oder lang sein.

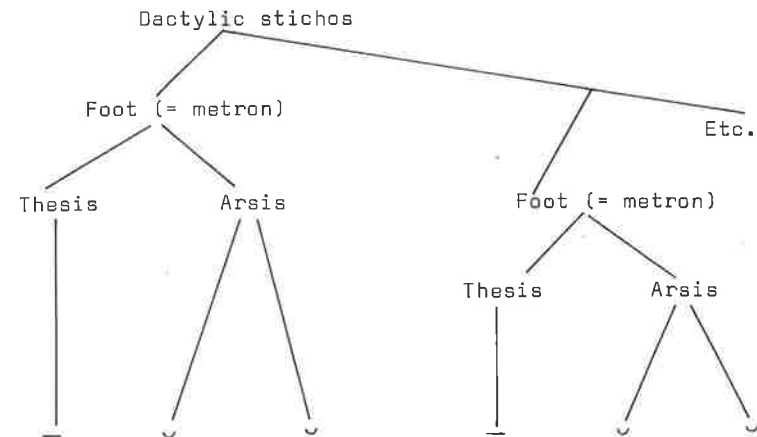
Für den ersten Vers der Aeneis z.B. ergibt sich entsprechend dieser Konzeption folgende Skandierung:

Arma virumque cano, Troiae qui primus ab oris

Das generative Modell von DEVINE/STEPHENS entspricht dagegen weitgehend der Konzeption der traditionellen Metrik.

Nimmt man einige für die weitere Argumentation unwichtige Vereinfachungen vor, lautet die Darstellung des klassischen (griechischen) Hexameters bei DEVINE/STEPHENS etwa folgendermaßen:

Mit Hilfe von metrischen Strukturregeln wird ein metrisches Basis-schema mit folgender Struktur erzeugt:



(DEVINE/STEPHENS 1975:413)

Dieses Schema wird mit Hilfe von Transformationsregeln in die verschiedenen metrischen Oberflächenschemata überführt:

1. SD - # v SC 1 3 2 optional (feminine caesura: further details omitted)
1 2 3

2. v → ∅ / _ # # (catalexis)

3. v → - optionally / _ # # (final anceps)

4. vv → - optionally / _ / 5 with rare exceptions (contraction)

(DEVINE/STEPHENS 1975:420)

13.2.2. In beiden Konzeptionen - das gleiche gilt auch für die Konzeption von BOWLEY (1974) - werden neben obligatorischen auch optionale (fakultative) Regeln verwendet.

Optionale Regeln tragen jedoch nicht der Tatsache Rechnung, daß auch die Anwendung dieser Regeln bestimmten Gesetzmäßigkeiten unterliegt, die allerdings nicht deterministischer, sondern probabilistischer Natur sind und eine probabilistisch verstandene Kompetenz widerspiegeln (vgl. CEDERGREN/SANKOFF 1974). Probabilistische Gesetzmäßigkeiten bei der Verwendung von Daktylen und Spondeen im lateinischen Hexameter können somit durch fakultative Zuordnungsregeln (HALLE) oder optionale Transformationsregeln (DEVINE/STEPHENS) nicht erfaßt werden.

Dies wird jedoch möglich, wenn man analog zu dem von SUPPES und anderen Forschern im Rahmen der Grammatiktheorie mit Erfolg praktizierten Verfahren (vgl. z.B. SUPPES 1972) Regeln probabilistisch bewertet (gewichtet), d.h. ihnen bestimmte Anwendungswahrscheinlichkeiten zuordnet.¹⁰

Optionale Regeln erhalten dabei eine Wahrscheinlichkeit aus dem offenen Intervall (0,1), die anhand der in einem bestimmten Korpus beobachteten Anwendungshäufigkeiten geschätzt wird. Obligatorischen Regeln wird im Fall der obligatorischen Anwendung die Wahrscheinlichkeit $p = 1$, im Fall der obligatorischen Nicht-Anwendung die Wahrscheinlichkeit $p = 0$ zugeordnet. Durch Multiplikation der Wahrscheinlichkeiten der einzelnen Regeln können dann die Wahrscheinlichkeiten der verschiedenen Ketten (Ableitungen) ermittelt werden. Bei der Bewertung der Regeln ist es möglich, nicht nur die Abhängigkeit von sprachlichen (kontextuellen), sondern auch von außersprachlichen Faktoren zu berücksichtigen. Auf diese Weise kann auch die bei CHOMSKY ausgesparte stilistisch-pragmatische Variation miterfaßt werden.

Im Rahmen einer solchen Konzeption kommt der in der generativen Linguistik meist vernachlässigten Korpusanalyse wieder eine wichtige Funktion zu. Im Gegensatz zum taxonomischen Strukturalismus ist die Korpusanalyse jedoch nicht das eigentliche Forschungsziel. Das eigentliche Ziel ist die Beschreibung und Erklärung von

Regularitäten im sprachlichen Verhalten, für das ein bestimmtes Korpus lediglich ein Beleg ist (vgl. KLEIN 1974a:54). Die Verwendung statistischer Verfahren ermöglicht dabei eine Projektion der Ergebnisse über das betreffende Korpus hinaus.

Anhand der Daten von DROBISCH (1866) soll nun im folgenden zuerst das Problem der Verteilung von Daktylen und Spondeen im lateinischen Hexameter genauer untersucht werden. Dabei werde ich auch auf die Arbeiten von BOLDRINI (1948), HERDAN (1954) und HERDAN (1966:206-209) eingehen, die ebenso wie WILLIAMS (1970:116-124) die Daten von DROBISCH erneut analysiert haben. Anschließend soll dann gezeigt werden, wie die gewonnenen Ergebnisse in ein generatives Modell des lateinischen Hexameters integriert werden können.

13.3. Ein probabilistisches Modell für den lateinischen Hexameter

13.3.1. Regressionsanalyse

13.3.1.1. Die von DROBISCH erhobenen Daten stammen aus der Aeneis und den Georgica von VERGIL und aus je einem Werk der Vergilischen Dichter LUKREZ, HORAZ, MANILIUS, PERSIUS, LUKAN, JUVENAL und der Nicht-Vergilischen Dichter ENNIUS, CICERO, CATULL, OVID, VALERIUS FLACCUS, STATIUS, SILIUS ITALICUS, CLAUDIAN. Der Untersuchungszeitraum erstreckt sich damit über etwa 600 Jahre (etwa 200 v. Chr. bis 400 n. Chr.).

DROBISCH entnahm jedem der ausgewählten Werke Blöcke von 80 Versen, solange bis ihm die relativen Häufigkeiten der einzelner Verstypen genügend stabil erschienen. Die in lateinischen Versen seltenen *spondiaci*¹¹, d.h. Verse mit einem Spondeus im fünften Fuß, blieben bei der Stichprobenentnahme unberücksichtigt. Der resultierende Stichprobenumfang betrug je nach Werk zwischen 414 und 649 Verse. Unter anderem hoffte DROBISCH auf diese Weise Unterschiede zwischen Vergilischen und Nicht-Vergilischen Dichtern hinsichtlich der Verteilung von Daktylen und Spondeen feststellen zu können.

Vernachlässigt man den seltenen *spondiacus* und die Auflösung der *syllaba anceps* am Versende in eine Länge oder in eine Kürze und berücksichtigt nur die ersten vier Versfüße, ergeben sich entsprechend der Anzahl der Variationen mit Wiederholung von Daktylen und Spondeen $V_4^W(2) = 2^4 = 16$ verschiedene Hexametertypen. Faßt man die bei DROBISCH (1866) für Vergil, für die Vergilischen Dichter, für die Nicht-Vergilischen Dichter und für alle Einzelstichproben gegebenen Daten zusammen, erhält man die in Tabelle 13.1 aufgeführten Häufigkeiten. Die 16 Verstypen sind in dieser Tabelle entsprechend der Anzahl der auftretenden Daktylen in 5 Gruppen eingeteilt und innerhalb der einzelnen Gruppen der Häufigkeit nach geordnet.

Die Tabelle 13.1 stimmt in einigen Punkten nicht mit den bei DROBISCH (1866), BOLDRINI (1948), HERDAN (1966) sowie WILLIAMS (1970) angeführten Daten überein. Während nämlich die Tabelle 13.1 auf den absoluten Häufigkeiten in den Einzelstichproben beruht, ist DROBISCH bei der Ermittlung der Daten für die Gesamtstichprobe und für die Dichtergruppen nicht von den absoluten Häufigkeiten ausgegangen, sondern hat Prozentwerte gemittelt, ohne Unterschiede im Stichprobenumfang zu berücksichtigen. Diese aufgrund des Vorgehens von DROBISCH nicht ganz korrekten Daten haben BOLDRINI, HERDAN und WILLIAMS dann offensichtlich übernommen, wobei WILLIAMS noch irrtümlicherweise den Stichprobenumfang bei VERGIL mit 2 Stichproben von jeweils 640 Versen und den Gesamtstichprobenumfang mit 9013 Versen angibt.

Tab. 13.1: Absolute und prozentuale Häufigkeiten des lateinischen Hexameters entsprechend der Anzahl und der Position von Daktylen (d) und Spondeen (s) nach DROBISCH (1866)

Nr.	Verstyp	Vergil f_{i1} %	Vergilianische Dichter (ohne Vergil) f_{i2} %	Nicht- Vergilianische Dichter f_{i3} %	Gesamt- Stichprobe f_i %
1	ssss	118 6,70	192 5,57	268 6,26	578 6,09
2	sssd	51 2,90	91 2,64	101 2,36	243 2,56
3	ssds	89 5,06	195 5,65	279 6,51	563 5,93
4	sdss	178 10,11	366 10,61	393 9,17	937 9,87
5	ds	263 14,94	529 15,34	620 14,47	1412 14,87
6	ssdd	32 1,82	70 2,03	74 1,73	176 1,85
7	sd	66 3,75	161 4,67	164 3,83	391 4,12
8	sdds	99 5,63	166 4,81	195 4,55	460 4,85
9	ds	110 6,25	240 6,96	267 6,23	617 6,50
10	dsds	199 11,31	394 11,42	535 12,49	1128 11,88
11	ds	208 11,82	363 11,10	471 10,99	1062 11,19
12	sddd	36 2,05	81 2,35	68 1,59	185 1,95
13	dsdd	65 3,69	131 3,80	198 4,62	394 4,15
14	ds	80 4,55	160 4,64	237 5,53	477 5,02
15	ds	126 7,27	221 6,41	299 6,98	646 6,83
16	ds	38 2,16	69 2,00	115 2,68	222 2,34
		1760 100,00	3449 100,00	4284 100,00	9493 100,00

13.3.1.2. Aus der Tabelle entnimmt man sofort, daß die Häufigkeiten der einzelnen Verstypen signifikant von einer Gleichverteilung abweichen. Der Erwartungswert für VERGIL unter der Nullhypothese beträgt $E_{i1} = 110$. Berechnen wir lediglich für O_{51} die Abweichung vom Erwartungswert, erhalten wir

$$\chi^2 = \sum_{i=1}^{16} \frac{(O_{i1} - E_{i1})^2}{E_{i1}} > \frac{(O_{51} - E_{i1})^2}{E_{i1}}$$

$$= \frac{(263 - 110)^2}{110} = 212,81 > \chi^2_{0,05,15} = 24,996.$$

In Tabelle 13.1 fällt auf, daß offensichtlich eine Tendenz besteht, in den vorderen Versfüßen vor allem Daktylen zu bevorzugen. Eine Zusammenfassung der entsprechenden Daten aus Tabelle 13.1 ergibt folgende Häufigkeiten:

Tab. 13.2: Absolute und prozentuale Häufigkeiten von Daktylen und Spondeen in den ersten vier Hexameterfüßen bei VERGIL (n = 1760)

Fuß	Daktylen		Spondeen	
1	1091	61,99	669	38,01
2	833	47,33	927	52,67
3	686	38,98	1074	61,02
4	478	27,16	1282	72,84
	3088		3952	

Der Trend kann mit Hilfe von Regressionsgleichungen beschrieben werden. Geht man von einem linearen Modell aus, lautet die allgemeine Form einer Regressionsgleichung (vgl. Kap. 12.2.2)

$$\hat{y} = a_{yx} + b_{yx}x$$

Die Parameter dieser Funktion lassen sich mit Hilfe der Methode der kleinsten Quadrate schätzen. Bezeichnen wir die abhängige Variable (Zielgröße) 'prozentuale Anzahl der Daktylen' mit y und die unabhängige Variable (Einflußgröße) 'Versfuß' mit x, erhalten wir für den Regressionskoeffizienten b_{yx} und den Achsenabschnitt a_{yx} folgende Werte:

$$b_{yx} = \frac{n \sum xy - \sum x \sum y}{n \sum x^2 - (\sum x)^2} = \frac{4(382,23) - 10(175,46)}{4(30) - 100} = -11,28$$

$$a_{yx} = \frac{\sum y - b_{yx} \sum x}{n} = \frac{175,46 + 11,28(10)}{4} = 72,08$$

Die Regressionsgerade lautet:

$$\hat{y} = 72,08 - 11,28x \quad (13.1)$$

Da zwischen der Anzahl der Daktylen und der Anzahl der Spondeen (w) die funktionale Beziehung $y = 100 - w$ besteht, kann die Regressionsgerade für w anhand der Regressionsgeraden für y berechnet werden:

$$\hat{w} = 27,92 + 11,28x \quad (13.2)$$

Setzen wir die x-Werte in die Gleichungen (13.1) und (13.2) ein, erhalten wir folgende Schätzwerte:

Tab. 13.3: Schätzwerte für die Häufigkeit von Daktylen und Spondeen (%)

Fuß	Daktylen	Spondeen
1	60,79	39,21
2	49,51	50,49
3	38,22	61,78
4	26,94	73,06

Ein Vergleich mit Tabelle 13.2 zeigt, daß die ermittelten Regressionsgeraden eine ziemlich genaue Vorhersage der Anzahl der Daktylen bzw. der Spondeen in den jeweiligen Versfüßen erlauben.

Die regressionsanalytische Darstellung ist somit in diesem Punkt erheblich adäquater als die Beschreibung durch optionale Regeln.

13.3.2. Die Ansätze BOLDRINI's und HERDAN's

BOLDRINI (1948) glaubte nachgewiesen zu haben, daß die Häufigkeiten der 5 Hexametergruppen hinlänglich gut ($\chi^2 = 2,605$; $P = 0,63$) durch das Bernoulligesetz beschrieben werden könnten, d.h. $B(4,p)$ verteilt seien. Mit $p_s = 0,558$ für den Spondeus und $p_d = 0,442$ für den Daktylus ergeben sich die theoretischen Werte nach der Formel:

$$f(x) = \binom{4}{x} 0,442^x 0,558^{4-x}, \quad x = 0, \dots, 4$$

Gleichzeitig stellte BOLDRINI jedoch fest, daß die ungruppierten Häufigkeiten der einzelnen Verstypen *n i c h t* dem Bernoulligesetz folgen ($\chi^2 = 28,953$; $P = 0,02$). Die entsprechenden Daten finden sich in der Tabelle 13.4.

Tab. 13.4: Beobachtete prozentuale Häufigkeiten ($\hat{p}_i$) und errechnete prozentuale Wahrscheinlichkeiten (p_i) der Hexametertypen nach BOLDRINI (1948) und HERDAN (1954)

Vers- typ	gruppiert		ungruppiert	
	$\hat{p}_i$	p_i	$\hat{p}_i$	p_i
ssss	6,10	9,69	6,1	9,69
sssd	33,61	30,72	2,6	7,68
ssds			6,0	7,68
sdss			10,0	7,68
dsss			15,01	7,68
ssdd	40,11	36,49	1,9	6,08
sdsd			4,11	6,08
sdds			4,8	6,08
dssd			6,5	6,08
dsds			11,8	6,08
ddss			11,0	6,08
sddd	17,78	19,28	2,0	4,82
dsdd			4,08	4,82
ddsd			5,0	4,82
ddds			6,7	4,82
dddd	2,40	3,82	2,4	3,82

Die Feststellungen BOLDRINI's veranlaßten HERDAN (1954) zu einer Überprüfung der von BOLDRINI durchgeführten Berechnungen. HERDAN stellte fest, daß BOLDRINI den χ^2 -Test nicht korrekt angewandt hat. Bei einer korrekten Anwendung dieses Tests ergeben sich folgende Werte:

gruppierte Häufigkeiten $\chi^2_4 = n (0,0260)$

ungruppierte Häufigkeiten $\chi^2_{15} = n (0,2895)$

Gehen wir von einem Signifikanzniveau von $P = 0,05$ aus, liegt der kritische Wert bei 9,49 bzw. bei 25,00. Dies hat zur Konsequenz, daß bereits bei $n = 364$ (gruppierte Daten) und $n = 86$ (ungruppierte Daten) eine signifikante Diskrepanz zwischen Theorie (Binomialmodell) und Beobachtung vorliegt.

Nachdem sich das Binomialmodell als inadäquat erwiesen hat, versuchte HERDAN (1954) das Problem mit Hilfe der Informationstheorie zu lösen. Er kam zu dem m.E. kaum gerechtfertigten Ergebnis, daß die Anzahl der Spondeen und Daktylen pro Hexameter nicht wesentlich von der nach dem Zufallsgesetz zu erwartenden abweicht, daß jedoch die Anordnung im wesentlichen auf den Einfluß der Wortstruktur¹² zurückzuführen sei.

Die Fragwürdigkeit dieses Ergebnisses zeigt sich bereits bei einer Überprüfung des Anteils von Spondeen ($p_s = 0,558$) und Daktylen ($p_d = 0,442$) in den ersten 4 Versfüßen. Unter der Nullhypothese beträgt der Erwartungswert der Daktylen $p = 0,5$. Die Abweichung der beobachteten Häufigkeit von diesem Wert ist jedoch hochsignifikant:

$$z = \frac{|\hat{p} - E(\hat{p})|}{\sqrt{V(\hat{p})}} = \frac{|\hat{p} - p|}{\sqrt{\frac{p(1-p)}{n}}} = \frac{|0,442 - 0,5| \sqrt{n}}{\sqrt{0,5(0,5)}} = 0,116 \sqrt{n}$$

Bei einem Signifikanzniveau von $P = 0,05$ wird der kritische Wert von $z = 1,96$ bereits bei $n = 285$ erreicht. Es besteht folglich trotz der erhöhten Häufigkeit von Daktylen am Versanfang eine Tendenz, in den ersten vier Versfüßen eher Spondeen zu verwenden.

Was die Anordnung der Daktylen und Spondeen betrifft, schlägt HERDAN (1966:209) selbst eine andere Interpretation vor. Da normalerweise der 5. Fuß im lateinischen Hexameter ein Daktylus ist,

tendiert der Dichter zur Vermeidung von Monotonie dazu, Daktylen eher am Versanfang zu gebrauchen. Diese Tendenz bestimmt dann auch - und hier argumentiert HERDAN genau entgegengesetzt zu seinen Ausführungen von 1954 - die Wortwahl. Die metrische Variation dient somit als Erklärung für die im Gegensatz zur Prosa freieren Wortstellung im lateinischen Hexameter.

HERDAN bleibt jedoch bei dieser m.E. wesentlich adäquateren Neuinterpretation des Phänomens stehen und macht keinen Vorschlag für ein aus diesen Annahmen abgeleitetes stochastisches Modell.

13.3.3. Das Modell einer nicht-homogenen MARKOV-Kette

13.3.3.1. Bei der Suche nach einem geeigneten stochastischen Modell muß zunächst geprüft werden, ob die Daten der einzelnen Dichtergruppen (vgl. Tab. 13.1) Zufallsvariablen gleicher Verteilung entstammen. Der Homogenitätstest mit $P = 0,05$ ergibt $\chi^2 = 47,22$. Da $\chi^2 = 47,22 > \chi^2_{30,0,05} = 43,773$ ist, ist die Hypothese der Homogenität zu verwerfen.

Dieses Ergebnis widerspricht dem von HERDAN (1954:293) festgestellten sehr hohen Grad von Übereinstimmung zwischen den Dichtern. Allerdings muß bei der Interpretation des Testergebnisses berücksichtigt werden, daß in die Berechnung ein relativ hoher Stichprobenumfang von $n = 9493$ eingegangen ist.

Außerdem sagt der χ^2 -Test auch nichts über die Stärke des Einflusses der Variablen 'Dichtergruppe' auf die Variable 'Verstyp' aus. Überprüft man die Stärke des Einflusses mit Hilfe des CRAMÉRSchen Kontingenzkoeffizienten (vgl. Kap. 12.2.5.1), ergibt sich lediglich ein äußerst schwacher Zusammenhang:

$$C_c = \sqrt{\frac{\chi^2}{n(k-1)}} = \sqrt{\frac{47,22}{9493(2)}} = 0,0499$$

Es ist deshalb durchaus gerechtfertigt, die drei Gruppen von Dichtern als homogen zu betrachten.

Die Homogenitätshypothese erscheint mir jedoch aus einem anderen Grund relativ problematisch. Man kann a priori nicht davon ausgehen, daß die verschiedenen Gruppen - es handelt sich immerhin um insgesamt 15 Dichter - in sich wiederum homogen sind. Es könnte vielmehr sein, daß auch die Variable 'Dichter' oder die Variable 'Zeit' einen Einfluß ausüben. Da auch die Ergebnisse von BOLDRINI (1948), der mit Hilfe der Varianzanalyse versucht hat, den Anteil der einzelnen Dichter an der Gesamtvarianz festzustellen, gegen die Homogenitätsannahme sprechen, beschränke ich mich im folgenden auf VERGIL.

13.3.3.2. Im Zusammenhang mit der Diskussion der Arbeit von BOLDRINI (1948) war gezeigt worden, daß die Annahme eines BERNOULLI-Modells inadäquat ist. Der Grund hierfür kann darin gesehen werden, daß bei diesem Modell die Unabhängigkeit der einzelnen Ereignisse vorausgesetzt wird, d.h. übertragen auf das vorliegende Problem, daß das Auftreten eines Daktylus bzw. eines Spondeus keinen Einfluß auf das Auftreten eines Daktylus bzw. eines Spondeus z.B. im folgenden Versfuß hat.

Wahrscheinlichkeitstheoretisch bedeutet dies, daß die Wahrscheinlichkeit einer Folge von n disjunkten Ereignissen das Produkt der Wahrscheinlichkeiten der Einzelereignisse ist, d.h. es gilt:

$$P(A_1, A_2, \dots, A_n) = P(A_1) P(A_2) \dots P(A_n)$$

Prüft man die Daten genauer, erweist sich die Unabhängigkeitsannahme jedoch als nicht gerechtfertigt. So scheint z.B. eine stark ausgeprägte Tendenz zu bestehen, auf Daktylen Spondeen folgen zu lassen oder mit anderen Worten, Daktylenfolgen zu vermeiden. Aufgrund der stark dissoziativen Wirkung des 5. (daktylischen) Versfußes scheint diese Tendenz zum 5. Versfuß hin zuzunehmen.

Wenn wir davon ausgehen, daß der Dichter bei der Wahl des i -ten Versfußes die Gestalt des $i-1$ -ten Fußes am stärksten berücksichtigt, erscheint das Modell einer nicht-homogenen MARKOV-Kette 1-ter Ordnung adäquat.¹³

MARKOV-Ketten¹⁴ 1-ter Ordnung sind diskrete stochastische Prozesse, bei denen die Wahrscheinlichkeit eines beliebigen Ereignisses E_j zum Zeitpunkt k - in der Terminologie der MARKOV-Ketten wird anstelle von Ereignissen auch von (System-)Zuständen gesprochen - von der Wahrscheinlichkeit eines beliebigen Ereignisses E_i zum Zeitpunkt $k-1$ bedingt ist, d.h. es gilt

$$P(E_j^{(k)} | E_i^{(k-1)}) = P_{ij}^{(k)}$$

In der Terminologie der MARKOV-Ketten heißt die bedingte Wahrscheinlichkeit $P_{ij}^{(k)}$ auch Übergangswahrscheinlichkeit vom Zustand E_i in den Zustand E_j zum Zeitpunkt k . Die Wahrscheinlichkeit des Anfangszustandes des Prozesses zum Zeitpunkt $k = t_0$ wird durch die unbedingte Wahrscheinlichkeit $P_i^{(k)}$ - die sog. Anfangswahrscheinlichkeit - gekennzeichnet.

Der Anwendung (in der Linguistik) liegt meist das (einfachere) Modell der sog. homogenen MARKOV-Ketten zugrunde, bei der die Wahrscheinlichkeit $P_{ij}^{(k)}$ nicht von dem Parameter k abhängt, also $P_{ij}^{(k)} = P_{ij}$ ist. Ein Beispiel hierfür sind Phonem- oder Graphemübergangswahrscheinlichkeiten, die in der Regel nicht vom Zeitpunkt k abhängen.

Die bisherigen Bemerkungen erlauben es nun, das vorliegende Problem in der Terminologie der MARKOV-Ketten zu formulieren. Betrachten wir die ersten vier Füße des Hexameters als eine Folge von Zufallsvariablen $(X_0, \dots, X_3)$, von denen jede den Wert $x_i \in \{d, s\}$ annehmen kann, dann bilden diese Zufallsvariablen eine MARKOV-Kette, wenn folgende Gleichung erfüllt ist:

$$P[(X_0, \dots, X_3) = (x_0, \dots, x_3)] = P(X_0 = x_0) \prod_{i=1}^3 P(X_i = x_i | X_{i-1} = x_{i-1}) \quad (13.3)$$

Die Verteilung der MARKOV-Kette $(X_0, \dots, X_3)$ ist also durch die Anfangsverteilung $\mu = (p_s^0, p_d^0)$ und durch die Übergangsmatrizen P_1, P_2 und P_3 festgelegt.

Die Berechnung der Übergangswahrscheinlichkeiten sei am folgenden Beispiel verdeutlicht.

Bezeichnen wir das Auftreten von Daktylen bzw. Spondeen in den ersten vier Versfüßen als $d_1, \dots, d_4$ bzw. als $s_1, \dots, s_4$, beträgt die Wahrscheinlichkeit, daß nach einem Daktylus im ersten Versfuß (d_1) ein Spondeus im zweiten Versfuß (s_2) auftritt, entsprechend der Definition der bedingten Wahrscheinlichkeit

$$P(s_2|d_1) = \frac{P(d_1 s_2)}{P(d_1)}.$$

Die Anfangswahrscheinlichkeit $P(d_1)$ und die Wahrscheinlichkeit der Folge $d_1 s_2$ schätzen wir anhand der beobachteten Häufigkeiten. Die Häufigkeit für d_1 erhält man unmittelbar aus Tabelle 13.2. Die Häufigkeit für $d_1 s_2$ ermittelt man aus Tabelle 13.1 durch Addition der Häufigkeiten der Verstypen 5, 9, 10 und 13. Mit $P(d_1) = 0,6199$ und $P(d_1 s_2) = 0,3619$ beträgt die gesuchte bedingte Wahrscheinlichkeit

$$P(s_2|d_1) = \frac{0,3619}{0,6199} = 0,5939.$$

Es ist üblich, die Übergangswahrscheinlichkeiten in Form von Matrizen zu schreiben. Wir erhalten für VERGIL folgende Übergangsmatrizen mit jeweils zwei Zuständen:

	d_2	s_2	
d_1	0,4161	0,5939	$= P_{ij}^{(1)} = P_1$
s_1	0,5665	0,4335	
	d_3	s_3	
d_2	0,3613	0,6387	$= P_{ij}^{(2)} = P_2$
s_2	0,4153	0,5847	
	d_4	s_4	
d_3	0,2493	0,7507	$= P_{ij}^{(3)} = P_3$
s_3	0,2858	0,7142	

Da die Ereignisse E_j paarweise disjunkt sind, gilt für die Zeilenvektoren dieser Matrizen $\sum_j P_{ij} = 1$.

Anhand der Übergangsmatrizen können nun die theoretischen Wahrscheinlichkeiten für die verschiedenen Hexametertypen berechnet werden. Z.B. ergibt sich für den Typ $s_1 s_2 s_3 s_4$ entsprechend der Definition der MARKOV-Kette:

$$P(s_1 s_2 s_3 s_4) = p_{s_1} P(s_2|s_1) P(s_3|s_2) P(s_4|s_3) \\ = 0,3601 (0,4353) 0,5847 (0,7142) = 0,0688$$

Die theoretischen Wahrscheinlichkeiten für alle 16 Hexametertypen sind zusammen mit den relativen Häufigkeiten in der Tabelle 13.5 aufgeführt.

Tab. 13.5: Relative Häufigkeiten $\hat{p}_i$ und errechnete Wahrscheinlichkeiten p_i von Hexametertypen bei VERGIL

Typ	$\hat{p}_i$	p_i
ssss	0,0670	0,0688
sssd	0,0290	0,0275
ssds	0,0506	0,0514
sdss	0,1011	0,0982
dsdd	0,1494	0,1537
ssdd	0,0182	0,0171
sdsd	0,0375	0,0393
sdds	0,0563	0,0584
dsdd	0,0625	0,0615
dsds	0,1131	0,1148
ddss	0,1182	0,1177

Fortsetzung von Tab. 13.5

Typ	$\hat{p}_i$	p_i
sddd	0,0205	0,0194
dsdd	0,0369	0,0381
ddsd	0,0455	0,0471
ddds	0,0727	0,0700
dddd	0,0216	0,0232

Wir überprüfen jetzt mit Hilfe des χ^2 -Tests, ob die Hypothese der MARKOV-Kette aufrechterhalten werden kann:

$$\chi^2 = \sum_{i=1}^{16} \frac{(n\hat{p}_i - np_i)^2}{np_i} = n \sum_{i=1}^{16} \frac{(\hat{p}_i - p_i)^2}{p_i} = 1760(0,0010) = 1,7399$$

Von den 14 Übergangswahrscheinlichkeiten mußten 7 Werte geschätzt werden. Die Anzahl der Freiheitsgrade beträgt deshalb $16 - 1 - 7 = 8$. Da $P(\chi_8^2 \geq 1,7399) \approx 0,99$ ist, kann mit einer statistischen Sicherheit von etwa 99 % die Hypothese der MARKOV-Kette beibehalten werden.

Die aus theoretischen Annahmen abgeleitete Hypothese der MARKOV-Eigenschaft ergibt somit ein im Verhältnis zur Binomialhypothese weit adäquateres Modell für den Vergilischen und vermutlich auch für den Hexameter der übrigen lateinischen Dichter. Möglicherweise läßt sich das Modell noch weiter verbessern, wenn man die Parameter mit Hilfe der Maximum-Likelihood-Methode schätzt.

Das MARKOV-Modell ist prädiktiv und erlaubt verschiedene Prognosen. Verwenden wir für das Prädikat "ist Vergilischer Hexameter" das Symbol H und für das Prädikat "gehört zum Hexametertyp ddss" das Symbol T_{11} , kann z.B. folgende Prognose formuliert werden (vgl. Tab. 13.5):

$$\begin{array}{lcl} G_1 : H(x) & \xrightarrow{p = 0,1177} & T_{11}(x) \\ A_1 : H(a) & & \\ \hline E_1 & & T_{11}(a) \end{array} \quad 0,1177$$

Dies bedeutet, daß beim Vorliegen der Antezedens-Bedingung A_1 aufgrund des 'Gesetzes' G_1 mit einer induktiven Wahrscheinlichkeit von 0,1177 auch das Explanandum E_1 gegeben ist.

Weiterhin können mit diesem Modell auch Fragen beantwortet werden wie z.B.: Wie groß ist die Wahrscheinlichkeit, daß im 4. Fuß ein Daktylus verwendet wird, wenn der Vers mit einem Daktylus beginnt?

Neben der weit größeren Übereinstimmung mit den Daten hat das MARKOV-Modell gegenüber dem Binomialmodell aber noch den weiteren Vorteil, daß es zumindest als ein erster Schritt zu einem komplexen Produktions- bzw. Perzeptionsmodell angesehen werden kann. Denn wie eine Vielzahl psychologischer Arbeiten gezeigt hat (vgl. z.B. LEE 1971), ist eine Interpretation menschlicher Tätigkeit als komplexer stochastischer Prozeß eine sinnvolle Hypothese zur Beschreibung und Erklärung der Struktur menschlichen Handelns.

13.3.3.3. Das MARKOV-Modell hat in der vorgeschlagenen Form den Nachteil, daß insgesamt 7 Parameter aus den Daten geschätzt werden müssen.

Es stellt sich nun die Frage, ob es nicht möglich ist, die für das MARKOV-Modell benötigten Parameter durch ein Regressionsmodell zu schätzen. Auf diese Weise könnte, sofern man ein Regressionsmodell mit weniger als 7 Parametern findet, die Anzahl der benötigten Parameter verringert werden.

In der folgenden Tabelle findet sich eine Zusammenstellung von 6 Übergangswahrscheinlichkeiten und einer Anfangswahrscheinlichkeit. Wie ein Blick in die Übergangsmatrizen $P_1, \dots, P_3$ zeigt, sind die in der Tabelle aufgeführten Übergangswahrscheinlichkeiten jeweils komplementär zu weiteren 6 Übergangswahrscheinlichkeiten. Das gleiche gilt für die Anfangswahrscheinlichkeit. Eine Vorhersage der 7 Tabellenwerte würde somit eine Vorhersage aller 14 Übergangswahrscheinlichkeiten und damit auch eine Vorhersage der Wahrscheinlichkeiten der einzelnen Hexametertypen erlauben.

Tab. 13.6: Anfangswahrscheinlichkeit eines Daktylus und Wahrscheinlichkeit des Übergangs von einem Daktylus zu einem Daktylus und von einem Spondeus zu einem Spondeus

Fuß	Daktylen	Spondeen
1	0,6199	
2	0,4161	0,4335
3	0,3613	0,5847
4	0,2493	0,7142

Sowohl bei den Daktylen als auch bei den Spondeen ist ein deutlicher Trend erkennbar, der in beiden Fällen das Modell einer Potenzfunktion der allgemeinen Form

$$y = ax^b$$

nahelegt.

Die Parameter dieser Funktion können wiederum mit der Methode der kleinsten Quadrate geschätzt werden. Dazu wird die Potenzfunktion mit Hilfe der Transformationen

$$y = \log y = Y \quad \text{und} \quad x = \log x = X$$

linearisiert. Wir erhalten die linearen Gleichungen

$$\log y = \log a + b \log x \quad \text{bzw.} \quad Y = \log a + bX.$$

Die zur Schätzung der Parameter notwendigen Rechnungen sind in der folgenden Tabelle aufgeführt.

Tab. 13.7: Schätzung der Regressionsparameter für die Werte der Daktylen (y) aus Tabelle 13.6

x	y	log y=Y	log x=X	XY
1	0,6199	-0,2077	0,0000	0,0000
2	0,4161	-0,3808	0,3010	-0,1146
3	0,3613	-0,4421	0,4771	-0,2110
4	0,2493	-0,6033	0,6021	-0,3632
		-1,6339	1,3802	-0,6888

Wir erhalten folgende Parameterwerte:

$$b = \frac{n \sum XY - \sum X \sum Y}{n \sum X^2 - (\sum X)^2} = \frac{4(-0,6888) + 1,3802(1,6339)}{4(0,6807) - (1,3802)^2} = -0,6113$$

$$\log a = \frac{\sum Y - b \sum X}{n} = \frac{-1,6339 + 0,6113(1,3802)}{4} = -0,1975$$

$$a = 0,6346$$

Die Koeffizienten der Regressionsgleichung für die Spondeen werden auf die gleiche Weise berechnet. Die beiden Regressionsgleichungen lauten:

$$y = 0,6346x^{-0,6113}, \quad x = 1, \dots, 4 \quad (13.4)$$

$$w = 0,2634x^{0,7215}, \quad x = 1, \dots, 3 \quad (13.5)$$

Durch Einsetzen der x-Werte in die Gleichungen (13.4) und (13.5) erhalten wir die Schätzungen für die Übergangswahrscheinlichkeiten aus Tabelle 13.6 und - aufgrund der Komplementarität - auch die Schätzungen für die übrigen Übergangswahrscheinlichkeiten:

$$P(d_1) = 0,6346$$

$$P(s_1) = 0,3654$$

	d_2	s_2	
d_1	0,4154	0,5846	= P'_1
s_1	0,5656	0,4344	

	d_3	s_3	
d_2	0,3242	0,6758	= P'_2
s_2	0,4181	0,5819	

	d_4	s_4	
d_3	0,2719	0,7281	= P'_3
s_3	0,2838	0,7162	

Ein Vergleich dieser Werte mit den direkt aus den Daten geschätzten Übergangswahrscheinlichkeiten zeigt, daß die ermittelten Potenzfunktionen eine ziemlich genaue Vorhersage der Übergangswahrscheinlichkeiten erlauben.

Die Wahrscheinlichkeiten der verschiedenen Hexametertypen erhalten wir wiederum durch Multiplikation der entsprechenden Übergangswahrscheinlichkeiten. Die ermittelten Werte sind in der folgenden Tabelle aufgeführt.

Tab. 13.8: Relative Häufigkeiten $\hat{p}_i$ und errechnete Wahrscheinlichkeiten p_i von Hexametertypen bei VERGIL (Regressionsmodell)

Typ	$\hat{p}_i$	p_i
ssss	0,0670	0,0662
sssd	0,0290	0,0262
ssds	0,0506	0,0477
sdss	0,1011	0,1000
dsss	0,1494	0,1546
ssdd	0,0182	0,0180
sdsd	0,0375	0,0396
sdds	0,0563	0,0488
dssd	0,0625	0,0613
dsds	0,1131	0,1129
ddss	0,1182	0,1276
sddd	0,0205	0,0182
dsdd	0,0369	0,0422
ddsd	0,0455	0,0506
ddds	0,0727	0,0622
dddd	0,0216	0,0232

Wir führen wiederum einen χ^2 -Test durch. Der Test ergibt diesmal einen Wert von $\chi^2 = 10,57$. Da wir jedoch lediglich die 4 Parameter der Regressionsgleichungen geschätzt haben, erhöht sich auch die Anzahl der Freiheitsgrade auf $16 - 1 - 4 = 11$. Die einem $\chi^2_{11} = 10,57$ entsprechende statistische Sicherheit beträgt etwa 50 %. Die anhand des Regressionsmodells geschätzten Übergangswahrscheinlichkeiten ermöglichen somit eine nicht so gute Anpassung wie die direkt aus den Daten geschätzten Werte (vgl. Tabelle 13.5). Möglicherweise lassen sich jedoch andere Regressionsgleichungen finden, die (bei gleicher Anzahl von Parametern) eine bessere Schätzung der Übergangswahrscheinlichkeiten ermöglichen. Die verwendete regressionsanalytische

Schätzung hat aber auf jeden Fall den Vorteil, daß lediglich zwei Funktionen mit jeweils zwei Parametern benötigt werden, um auf der Basis des MARKOV-Modells die Auftretenswahrscheinlichkeiten aller 16 Hexametertypen vorherzusagen.

13.4. Probabilistische generative Metrik

Es gibt eine Reihe von Möglichkeiten, die ermittelten probabilistischen Gesetzmäßigkeiten in ein generatives Modell zu integrieren. Es könnten z.B. die Korrespondenzregeln von HALLE (1970), die metrischen Strukturregeln von BOWLEY (1974) oder die Transformationsregeln von DEVINE/STEPHENS (1975) probabilistisch bewertet werden.¹⁵ Welchen Weg man wählt, hängt letztendlich davon ab, was für ein Ziel man mit einer generativen Beschreibung verfolgt und welche - wiederum von der jeweiligen Zielsetzung abhängenden - Adäquatheitsbedingungen (z.B. formale Einfachheit vs. Natürlichkeit) man für generative Modelle akzeptiert (vgl. VON STECHOW 1970).¹⁶

Verzichtet man auf Transformationsregeln, könnten die verschiedenen Hexametertypen z.B. durch folgende metrische Strukturregeln (Ersetzungsregeln) erzeugt werden:

$$R_0: H \rightarrow F_j, \quad j = 1, \dots, 6$$

$$* R_{ij}: F_j \rightarrow A_i, \quad i = 1, 2, 3$$

$$A_1 = d, \quad A_2 = s, \quad A_3 = t$$

Die Regel R_0 bedeutet, daß ein Hexameter (H) aus 6 Versfüßen besteht. Die Regel R_{ij} besagt, daß ein Versfuß je nach Anwendungsfall (i, j) entweder einem Daktylus (d), einem Spondeus (s) oder einem Trochäus (t) entspricht.

Für die Anwendung der Regel R_0 gilt $P(R_0) = 1$.

Die Anwendung der Regel R_{ij} geschieht mit den Anfangswahrscheinlichkeiten $P(R_{i1})$ und den Übergangswahrscheinlichkeiten $P(R_{ij}|R_{ij-1})$ für $j = 2, \dots, 5$. Außerdem gilt:¹⁷

$$P(R_{3j}) = 0, \quad j = 1, \dots, 5$$

$$P(R_{1j}) = 0, \quad j = 6$$

Für $i = 1, 2$ und $j = 1, \dots, 4$ sind die Bewertungen durch die Potenzfunktionen (13.4) und (13.5) bzw. durch die Anfangsverteilung μ und die Übergangsmatrizen P_1, P_2 und P_3 gegeben. Ob auch für $j = 5, 6$ bedingte Wahrscheinlichkeiten anzunehmen sind, kann wegen fehlender Daten nicht entschieden werden.

Um deutlich zu machen, daß die Übergangswahrscheinlichkeiten die Wahrscheinlichkeit des Übergangs von Regel R_{ij} zu Regel R_{ij+1} angeben,¹⁸ könnten die Übergangsmatrizen in folgender Form geschrieben werden:

	R_{12}	R_{22}	
R_{11}	0,4161	0,5939	$= P(R_{i2} R_{i1}) = P_1$
R_{21}	0,6565	0,4335	

Die Wahrscheinlichkeit der Verwendung eines bestimmten Hexametertyps (Kette) wird jetzt als eine Funktion der Anwendungswahrscheinlichkeit der Regeln R_0 und R_{ij} betrachtet. Die Werte dieser Funktion werden durch Multiplikation der entsprechenden Anwendungswahrscheinlichkeiten ermittelt. Die Regel R_0 kann bei der Berechnung unberücksichtigt bleiben, da $P(R_0) = 1$ ist.

Dieses Modell kann auf verschiedene Weise erweitert werden. Es könnten z.B. Zäsurregeln eingeführt werden. Außerdem könnten auch die Zuordnungsregeln probabilistisch bewertet werden. Auf diese Weise würde z.B. eine eventuelle Tendenz zur Verwendung bestimmter Silben-

typen oder auch von Monosyllaba in bestimmten Verspositionen Berücksichtigung finden.

Von besonderem Interesse scheint mir eine Erweiterung in Richtung auf eine metrische Varietätengrammatik zu sein, die analog zu den von KLEIN (1974a:9) genannten Zielsetzungen einer Varietätenlinguistik

- (1) jede einzelne metrische Varietät beschreibt,
- (2) den Zusammenhang zwischen den einzelnen Varietäten darstellt und
- (3) die Varietäten zu außersprachlichen Faktoren in Beziehung setzt.

In bezug auf den lateinischen Hexameter könnte dabei folgendermaßen vorgegangen werden: Mit Hilfe von außersprachlichen Faktoren wie Autor, Genre oder Zeit werden verschiedene Hexameter-varietäten (z.B. Aeneis des VERGIL, Annalen des ENNIUS) definiert. Für diese Varietäten wird eine übergeordnete 'Bezugsgrammatik' formuliert, deren Regeln dann für jede einzelne Varietät probabilistisch bewertet werden. Auf diese Weise könnten sowohl Gemeinsamkeiten als auch autor-, genre- oder zeitspezifische Besonderheiten lateinischer Hexametervarietäten erfaßt werden.

Auch interlinguale typologische Vergleiche - z.B. zwischen dem lateinischen und dem griechischen Hexameter - könnten auf diese Weise durchgeführt werden. Tertium comparationis wäre dabei die übergeordnete Bezugsgrammatik. Verglichen würde die Anwendung der Regeln dieser Grammatik in den Vergleichssprachen.

14. ZUSAMMENFASSUNG UND AUSBLICK

In dieser Arbeit sollte aufgezeigt werden, wie auf der Basis linguistischer und statistischer Modellbildung ein Beitrag zu einer stärkeren 'Objektivierung' und 'Empirisierung' der Metrik geleistet werden kann. Dazu erschien es mir nötig, zuerst einmal auf der Grundlage der Analytischen Wissenschaftstheorie und speziell des Kritischen Rationalismus das von mir zugrundegelegte Verständnis von empirischer Wissenschaft zu explizieren und daran anschließend ein Rahmenmodell für eine empirisch-kommunikationswissenschaftlich orientierte Metriktheorie zu konzipieren. Da ein wichtiger Teil erfahrungswissenschaftlicher Theoriebildung - verstanden als konzeptuelle Rekonstruktion der Realität - in der Konstruktion präziser und theoretisch fruchtbarer Begriffe zu sehen ist, werden in den Kapiteln 4, 5, 7 und 8 einige für die Metrik wichtige Begriffe wie Text, Zufallsfolge, Rhythmus, Metrum, Vers, Silbe, Reim einer genaueren Analyse unterzogen. Ein weiterer Schwerpunkt dieses ersten Teils der Arbeit liegt in der Untersuchung der Interdependenz von metrischen und sprachlichen Strukturen (Kap. 6, 8.5 und 9). Dabei wird die entscheidende Bedeutung des Explanans 'Sprache' für die Analyse metrischer Strukturen aufgezeigt und damit gleichzeitig auch eine 'ontologische' Begründung für die Notwendigkeit einer verstärkten Einbeziehung der Linguistik in die Theoriebildung der Metrik gegeben.

Der zweite Teil der Arbeit beginnt mit einer Diskussion des Beitrags mathematischer Methoden und Modelle für die Entwicklung einer empirischen Metriktheorie (Kap. 11). Es wird gezeigt, daß in der Metrik, wie auch in allen anderen empirischen Wissenschaften, erst durch die Verwendung mathematischer Methoden eine tiefere Einsicht in den Gegenstandsbereich ermöglicht wird und daß es deshalb die Aufgabe der Metrik ist, "das zu messen, was meßbar ist, und das, was nicht meßbar ist, meßbar zu machen" (GALILEI). In diesem Zusammenhang wird vor allem die Bedeutung quantitativer und insbesondere stochastischer Methoden für die Textanalyse hervorgehoben. Denn diese Methoden erlauben, neben den manifesten, kategorischen Eigen-

schaften auch latente Tendenzen zu erfassen. Dadurch wird auch eine adäquate Analyse der stilistisch-pragmatisch bedingten metrischen Variation ermöglicht. Im Anschluß an die Bemerkungen zur Notwendigkeit einer verstärkten Mathematisierung der Metrik wird ein Begriffsrahmen für die quantitative Textanalyse skizziert, der forschungslogische Ablauf quantitativer Textanalysen dargestellt und das Problem der Skalierung und Messung textueller Eigenschaften und Relationen diskutiert.

Die Ausführungen des Kapitels 11 sind u.a. als theoretische Fundierung der sich im Kapitel 12 anschließenden quantitativen Analyse metrischer Texte intendiert. Diese hat primär das Ziel, anhand eines beschränkten Korpus einige Möglichkeiten der quantitativen Analyse formaler Textstrukturen aufzuzeigen. Dabei wird auch anhand eines Beispiels gezeigt, wie mit Hilfe empirischer Rezeptionsanalysen die Textanalyse weiter objektiviert werden kann. Im Rahmen des Kapitels 12 werden u.a. folgende statistische Modelle verwendet: Regressions- und Korrelationsrechnung, Iterationstheorie, Varianzanalyse und eine Reihe weiterer parametrischer und nicht-parametrischer Testverfahren. Dabei werden soweit wie möglich auch die mathematisch-statistischen Grundlagen der verwendeten Modelle dargestellt.

Zum Abschluß der Arbeit (Kap. 13) wird die jüngste Entwicklung in der Metrik, die generative Metrik, diskutiert und anhand des lateinischen Hexameters nachgewiesen, daß diese Richtung wichtige Aspekte metrischer Texte nicht adäquat zu analysieren vermag. Der deterministisch verstandenen Kompetenz der generativen Metrik wird ein probabilistischer Kompetenzbegriff entgegengesetzt und gezeigt, wie die generative Metrik durch eine aus einem MARKOV-Modell abgeleitete probabilistische Regelbewertung und durch Einbeziehung regressionsanalytischer Modelle zu einer erfahrungswissenschaftlichen Theorie mit erklärender und prädiktiver Kraft weiterentwickelt werden kann.

Im Rahmen dieser Arbeit konnte nur ein Teil der bei der Entwicklung einer empirischen, dem Postulat der intersubjektiven Überprüfbarkeit genügenden Metriktheorie sich stellenden Probleme diskutiert werden. Zukünftige, die skizzierte Konzeption weiterführende

Arbeiten sollten aufgrund der Komplexität des Gegenstandsbereiches und der daraus resultierenden Notwendigkeit der Verwendung immer komplexerer Methoden soweit wie möglich in Zusammenarbeit vor allem von Linguisten, Literaturwissenschaftlern, Mathematikern (Statistikern), Psychologen und Sozialwissenschaftlern realisiert werden. Zur Prüfung der theoretischen Annahmen und für die Heuristik der Hypothesenformulierung müßten eine Vielzahl verschiedener Sprachen und Texte untersucht werden. Dabei dürfte der Mathematik nicht nur als Methode der Sprach- und Textanalyse, sondern auch als gemeinsame, interdisziplinäre Metasprache eine besondere Bedeutung zukommen.

ANMERKUNGEN

1. Metrik und Linguistik

- 1 Neben dem Terminus 'Metrik' werden - z.T. synonym, z.T. auch in unterschiedlicher Bedeutung - Termini wie Verstheorie, Versforschung, Verskunst oder Verslehre verwendet. In englischsprachigen Publikationen finden sich vor allem die Termini *metrics* und *versification*. Ich gehe von einem weiten Verständnis von Metrik aus, das die nicht-normativen Verwendungsweisen der genannten Termini weitgehend mit einschließt.

- 2 IHWE (1975) verwendet den Terminus "metrically ordered text", der zwar das Spezifische dieser Textsorte genauer bezeichnet, allerdings auch weniger handlich ist.

- 3 Bereits ARISTOTELES kritisierte im ersten Kapitel seiner "Poetik" die *ars metrica*:

πλὴν οἱ ἄνθρωποι γε συνάκτοντες τῷ μέτρῳ τὸ ποιεῖν ἐλεγείῳ ποιοῦς τοὺς δὲ ἐποικοῦς ὀνομάζουσιν, οὐχ ὡς κατὰ τὴν μῆσιν ποιητὰς ἀλλὰ κοινῇ κατὰ τὸ μέτρον προσαγορεύοντες· καὶ γὰρ ἂν ἱατρικὸν ἢ φυσικὸν τι διὰ τῶν μέτρων ἐκφέρωσιν, οὕτω καλεῖν εἰσάθουσιν· οὐδὲν δὲ κοινὸν ἐστὶν Ὀμήρῳ καὶ Ἐμπεδοκλεῖ πλὴν τὸ μέτρον, διὸ τὸν μὲν ποιητὴν δόκαλον καλεῖν, τὸν δὲ φυσικὸν λόγον μᾶλλον ἢ ποιητὴν.

[Kassel (1965), *Aristotelis de arte poetica liber*, 1447b)
"In Wirklichkeit verbinden die Leute das Dichten mit dem Vers-
machen und nennen die einen Elegiker, die andern Epiker, nicht
im Hinblick auf die Art der Nachahmung, sondern ganz allgemein
nach dem Versmaße. Selbst wenn einer medizinische oder natur-
wissenschaftliche Dinge in Versen vorbringt, pflegt man so zu
reden. Homer und Empedokles haben indessen nichts gemeinsames
außer dem Versmaße. So wäre es richtig, den einen Dichter zu
nennen, den anderen aber eher Naturforscher." Übersetzung:
GIGON (1950:392)

- 4 Die Aufzählung berücksichtigt nicht die Tatsache, daß Stilistik und Poetik wiederum zur Linguistik bzw. Literaturwissenschaft gerechnet werden können.

- 5 Seit der griechisch-römischen Antike gilt die Metrik in erster Linie als ein Teilgebiet der Poetik und noch 1925 wird diese Auffassung von dem bekannten russischen Metriker ŽIRMUNSKIJ (1966:17) vertreten.

- 6 Das im folgenden angesprochene Verhältnis der Linguistik zur Literaturwissenschaft bzw. zur Poetik ist seit einigen Jahren äußerst umstritten. Aus der umfangreichen Literatur zu dieser Problematik sei lediglich auf die grundlegende Dissertation von IHWE (1972), auf KLOEPFER (1975) und auf GROEBEN (1976) hingewiesen.

- 7 IHWE (1976) sieht die Linguistik als "vorangehende Disziplin" (TARSKI) der Literaturwissenschaft und weiterer sprachthematizierender Disziplinen an. Diese Disziplinen sind als abhängig vom Stand und den Annahmen der Linguistik zu betrachten.

- 8 Ähnliches gilt auch für einige weitere "Text-Disziplinen" wie Philosophie oder Theologie. So gibt es z.B. seit einiger Zeit eine interdisziplinäre Zeitschrift für Theologie und Linguistik, die "Linguistica Biblica", in der u.a. Artikel zur "Generativen Poetik des Neuen Testaments" erschienen sind.

- 9 Es muß allerdings betont werden, daß der sprach-orientierte Zugang zur Metrik keineswegs neu ist. Schon seit der griechisch-römischen Antike haben sich immer wieder Grammatiker zu Fragen der Metrik geäußert.

- 10 Vgl. die Übersicht bei TARANOWSKI (1963). Einen Eindruck der weiteren Entwicklung vermitteln auch die Kongreßberichte der Metrikongresse von 1964 und 1966 in Brno (vgl. *Teorie Verše* 1966, 1968) und die beiden Bände *Poetics*, *Poetyka*, *Poëtika* (1961, 1966). Weitere Hinweise zur Entwicklung bis 1971 finden sich bei ČERVENKA (1971, Anm. 1).

- 11 Vgl. z.B. Heft 3 (1971) der "Zeitschrift für Literaturwissenschaft und Linguistik" oder die Hefte 12 (1974) und 16 (1975) der Zeitschrift "Poetics", die ausschließlich der Metrik gewidmet sind. Insgesamt gesehen ist die Anzahl der Veröffentlichungen zur Metrik mittlerweile unüberschaubar geworden. Die Bibliographie von THIEME (1933) verzeichnet z.B. allein für das Französische bis zum Jahre 1932 über 1000 Titel und die Literaturberichte von KALINKA (1935, 1937) etwa 800 (kommentierte) Titel für das Lateinische. Auch die Literaturberichte von CLARKE (1937) für das Spanische, von SHAPIRO (1948) für das Englische, von ŠTOKMAR (1933) für das Russische und von SCHEEL (1967, 1969) für die Romania - um nur einige Titel zu nennen - geben ein Beispiel für die Fülle (traditioneller) Veröffentlichungen. Diese Tatsache scheint jedoch einer beträchtlichen Anzahl von Metrikern nicht bekannt zu sein. Als Beispiel für eine explizite Fehleinschätzung der Literaturlage sei AVALLE (1963:8) zitiert, der meint, daß die Forschung das Interesse an der Metrik verloren habe: "Ora, al giorno d'oggi, non c'è quasi nessuno che se ne occupi." Der Literaturbericht von SCHEEL zeigt jedoch, daß eher folgendes Urteil angemessen ist: "Wenn man die Fülle neuer Beiträge überblickt, die auf dem Gebiet der Versforschung in den letzten anderthalb Jahrzehnten erschienen sind, muß man sagen, daß die Wissenschaft vom Vers - innerhalb und außerhalb der Romania - noch nie eine solche Aktivität an den Tag gelegt hat, wie gerade jetzt." (SCHEEL 1967:40). Zur Einschätzung der Literaturlage vgl. auch FAURE (1970:15).

- 12 Der Terminus 'strukturalistische Metrik' wird von mir in Abgrenzung zur generativen Metrik gebraucht und umfaßt z.B. sowohl am (nordamerikanischen) taxonomischen Strukturalismus als auch am (tschechischen) Funktionalismus orientierte Richtungen. Vgl. zur Definition von strukturalistischer Linguistik z.B. BIERWISCH (1966:78), LEPSCHY (1969:16) und FIGGE (1973). 'Mathematische Metrik' soll global alle die Richtungen bezeichnen, die bei der Analyse metrischer Texte *systematisch* mathematische Modelle verwenden. Die Tatsache, daß der Strukturbegriff auch für die generative Grammatik grundlegend ist und daß die generative Grammatik auch als Teildisziplin der mathematischen Linguistik aufgefaßt werden kann, wird bei dieser terminologischen Abgrenzung vernachlässigt.
- 13 Vgl. z.B. MENNEMEIER (1971:197): "Das Problem des freien Rhythmus darf nicht allein im engen Rahmen einer Metrik und schon gar nicht in der Weise normativer Absolutsetzung einer bestimmten geschichtlichen Erscheinung behandelt werden. Jenes Problem hat seine Stelle vielmehr in einem umfassenden poetologischen, literaturgeschichtlichen und selbst ideengeschichtlichen Kontext, von dem her erst der ganze Reichtum des betreffenden Gegenstandes *hermeneutisch* (Sperrung R.G.) erschlossen werden kann. In diesem Sinne sollen mögliche Aspekte des Themas 'Freier Rhythmus' im folgenden erörtert werden:
- 1) Klopstocks (freier) Rhythmus, verstanden als eine subjektive rhythmische Ausprägung der in der damaligen Epoche virulent werdenden Idee der Freiheit;
 - 2) die transzendentalphilosophische Konzeption eines freien Rhythmus in der deutschen Frühromantik,..."
- Auf Betrachtungen dieser Art werde ich verzichten.
- 14 Eine gemeinsame (vergleichende) Darstellung dieser vier Richtungen existiert m.W. bisher noch nicht.

2. Zur wissenschaftstheoretischen Basis einer empirischen Metrik
- 1 Einen Überblick über die vielfältigen Strömungen innerhalb der Analytischen Wissenschaftstheorie vermittelt das grundlegende Werk von STEGMÜLLER 1969-1974
 - 2 Im Rahmen der von mir vertretenen Konzeption einer Metriktheorie kann m.E. der *non-statement view* wissenschaftlicher Theorien (vgl. STEGMÜLLER 1973: z.B. 191) vernachlässigt werden. Es bleibt zu erwägen, ob evtl. die generative Metrik im Rahmen eines *non-statement view* zu rekonstruieren ist, wie es KANNGIESSER (1976) für formale Grammatiken vorschlägt. Vgl. hierzu auch TER MEULEN (1976).
 - 3 Sehr deutlich ist dies von KUHN (1970) anhand von Beispielen aus der Wissenschaftsgeschichte belegt worden. KUHN zeigt, daß ein Paradigmenwechsel gleichzeitig zu einem "Gestaltwechsel" führt, d.h. zu einer Uminterpretation von Beobachtungsfakten. Die Sprach- und Theorieabhängigkeit gilt jedoch nicht nur für die wissenschaftliche Erkenntnis, sondern auch, wie die Wahrnehmungspsychologie gezeigt hat, für jegliche Art der Erkenntnis (vgl. HOLZKAMP 1973). In der Linguistik wird diese Problematik im Zusammenhang mit der SAPIR-WHORF Hypothese diskutiert.
 - 4 Ansätze, den Wahrheitsbegriff auf ein Konsenskriterium zurückzuführen, finden sich bei POPPER bei der Entscheidung über die Anerkennung eines Basissatzes als Falsifikator. Zur Begründung einer Konsenstheorie unter weitgehendem Verzicht auf ein semantisches Wahrheitskriterium vgl. HABERMAS (1973).
 - 5 Vgl. POPPER (1973:10ff). Zu den verschiedenen Fassungen des Sinnkriteriums bei CARNAP und WITTGENSTEIN vgl. STEGMÜLLER (1965:380ff und 465ff).
 - 6 CARNAP schlägt deshalb vor, anstelle der Verifizierbarkeit bzw. der Falsifizierbarkeit lediglich die Bestätigungsfähigkeit bzw. Prüfbarkeit von Aussagen zu verlangen. Vgl. zur Bedeutung dieser beiden Begriffe STEGMÜLLER (1965:403f).
 - 7 Diese Begriffe gehen auf REICHENBACH (1938:6ff) zurück, der zwischen "context of discovery" und "context of justification" unterscheidet. In der neueren Diskussion wird neben dem Entdeckungszusammenhang als weitere pragmatische Dimension wissenschaftlicher Aussagensysteme der Wirkungs- bzw. Verwendungszusammenhang unterschieden (vgl. FRIEDRICH 1973:54). Die Vernachlässigung des Entstehungs- und Wirkungszusammenhangs durch den Kritischen Rationalismus ist von Gegnern der Analytischen Wissenschaftstheorie (z.B. von der Kritischen Theorie) vor allem unter gesellschaftskritischen Gesichtspunkten (Relevanz wissenschaftlicher Forschung) heftig kritisiert worden (vgl. z.B. ADORNO u.a. 1974). Allerdings ist die Kritik m.E. z.T. unbegründet. Vgl. die Stellungnahme zu diesem Problem bei ALBERT (1968:37ff).

- 8 So meint z.B. RUDNER (1966:6) es gäbe keine "logic of discovery". Diese Ansicht dürfte doch, wie die Arbeit von SEREBRJANNIKOV (1975) zeigt, als widerlegt gelten. Dies spricht dafür, auch den Entdeckungszusammenhang zum Objektbereich eines normativ verstandenen Kritischen Rationalismus zu zählen.
- 9 POPPER verlangt nicht, daß empirische Aussagen überprüft werden müssen. Eine solche Forderung würde zu einem infiniten Regreß führen (vgl. POPPER 1973:21).
- 10 FEYERABEND (1975) zieht aus der Unsicherheit der empirischen Basis die Konsequenz, sich in den Skeptizismus zu flüchten.
- 11 Der Entwicklung solcher Theorien wird in der empirischen Sozialforschung in jüngster Zeit - z.B. im Zusammenhang mit der Befragung als Datenerhebungsinstrument - eine besondere Bedeutung beigemessen.

3. Ein Rahmenmodell zur Analyse metrischer Texte

- 1 Diese Tatsache wird erst seit relativ kurzer Zeit von der Linguistik stärker berücksichtigt. Im Anschluß an F. DE SAUSSURE wurde vor allem der Systemcharakter von Sprache untersucht, das Funktionieren von Sprache in konkreten Kommunikationsprozessen und die Textqualität sprachlicher Äußerungen jedoch weitgehend vernachlässigt. Der Tätigkeitscharakter von Sprache ist vor allem von LEONT'EV (1971) herausgearbeitet worden.
- 2 Vgl. z.B. SCHMIDT (1973,1975); WIENOLD (1971,1972); IHWE (1972); GROEBEN (1972). SCHMIDT (1973) intendiert auch eine Neuorientierung der Linguistik und betrachtet ähnlich wie IHWE eine solche Texttheorie als Grundlagentheorie jeglicher Art verbaler Interaktion und damit als Basis für Soziolinguistik, Psycholinguistik, Literaturwissenschaft usw. WIENOLD vertritt einen allgemeinen semiotischen Ansatz. GROEBEN betont als Psychologe vor allem den literaturpsychologischen Aspekt des Rezipientenverhaltens.
- 3 Die Pragmatik spielt in den einzelnen textwissenschaftlichen Modellen eine unterschiedliche Rolle. Während z.B. VAN DIJK (1972) seiner Textgrammatik lediglich eine pragmatische Komponente additiv hinzufügt (vgl. die Kritik bei RIESER/WIRRER 1974:75 f), ist z.B. bei SCHMIDT (1973) oder bei BREUER (1974) die Pragmatik ein integrativer Bestandteil.
- 4 Dieses mit Absicht stark simplifizierende Modell soll nur einige für die unmittelbar folgende Argumentation wichtige Aspekte verdeutlichen. Vgl. zum "Verkürzungsmerkmal" als allgemeines Merkmal von Modellen STACHOWIAK (1965:437ff).
- 5 Vgl. GROEBEN (1976:125), der die Situation in der traditionellen Literaturwissenschaft folgendermaßen beurteilt: "Als Hauptgrund für das Verfehlen einer wissenschaftlichen Objektivität (im Sinne intersubjektiver Nachprüfbarkeit/Falsifizierbarkeit) kann man die Subjekt-Objekt-Konfundierung der hermeneutischen Interpretationsverfahren ansehen." Eine ausführliche kommunikationswissenschaftliche Analyse der textwissenschaftlichen Forschungssituation findet sich bei RIEGER (1972).
- 6 Vgl. WIMSATT/BEARDSLEY (1964:193): "A poem, as verbal artifact or complex linguistic entity, is, to be sure, actualized and realized in particular performances of it - in being read silently or aloud. But the poem itself is not to be identified with any performance of it or with any subclass of performances."
- 7 Vgl. WELLS (1964:197ff), der zwischen "dem Gedicht" (*the poem*) und den schriftlichen Fassungen des Gedichts (Manuskript des Autors und darauf beruhende schriftliche Fassungen) unterscheidet. "Das Gedicht" ist für den Interpreten lediglich anhand einer schriftlichen Fassung erschließbar, die allerdings aufgrund der ungenügenden Abbildfunktion der Schrift nur eine partielle Repräsentation der Ausdrucksseite "des Gedichts" darstellt.

- 8 Die Kontroverse zwischen Nominalisten und Essentialisten (Realisten) hat philosophiegeschichtlich eine lange Tradition. Bekannt ist in diesem Kontext vor allem der mittelalterliche Universalienstreit.
- 9 IHWE (1975) sieht eine Möglichkeit zur Lösung dieses Problems in der vollen Anwendung der phonologischen Komponente der generativen Transformationsgrammatik auf die Analyse metrischer Texte.

4. Texte und Zufallsfolgen

- 1 Eine umfassende Anwendung der Theorie der Binärfolgen auf die Metrik stellt m.E. ein - vermutlich nur durch die Zusammenarbeit von Metrikern und Mathematikern realisierbares - Forschungsdesiderat dar.
- 2 Die Menge aller aus dem Grundalphabet M gebildeten Folgen zusammen mit der Addition ist eine freie Halbgruppe über M. Anstelle von Addition wird im Zusammenhang mit freien Halbgruppen (freien Monoiden) meist von Verkettung gesprochen.
- 3 Z.B. ist im Deutschen eine unmittelbare Nachbarschaft zweier hauptbetonter Silben innerhalb eines Wortes nicht möglich.
- 4 IHWE (1975:380) spricht von "supplementary ordering principles" und begründet das Prädikat '*supplementary*' folgendermaßen: "The primary functions of human communication via natural languages are not disturbed when a natural language exhibits only rudimentary forms of metrical systems, or none at all. (...)there is no minimum of metrical system that a natural language must manifest in order to function and continue to function."
- 5 Vgl. auch KOCH (1971:374): "Ein Telefonbuch, ein Katalog, ein Buch, das nur aus Illustrationen von Atommodellen besteht, weisen zahlreiche Rekurrenzstrukturen auf, welche die normale Metrik weit übertreffen."

5. Rhythmus und Metrum

- 1 Vgl. CHATMAN (1965:12): "In too many discussions of meter one finds confusion between rhythm and meter." Auch in der Musikwissenschaft werden Rhythmus und Metrum nicht eindeutig verwendet. Vgl. hierzu die Feststellung des Musikwissenschaftlers BENARY (1967:8): "Trotzdem steht eine verbindliche terminologische Abgrenzung der Metrik von der Rhythmik immer noch aus." (Vgl. auch WINCKEL (1960:75f).
- 2 Einen Überblick vermittelt die grundlegende Arbeit von FRAISSE (1956). Dort findet sich auch eine Übersicht über verschiedene Rhythmusdefinitionen. Wichtige ältere Arbeiten sind u.a.: BOLTON (1894); MEUMANN (1894,1896); WALLIN (1901); BENUSSI (1913); SONNENSCHNEIN (1925); SERVIEN (1930); DE GROOT (1932).
- 3 Die vermutliche erste Definition von Rhythmus im heutigen Sinne findet sich bei PLATO, Gesetze II 665a: "τῆ δὲ τῆς κλυήσεως τάξεως ὁνομαζόμενα εἴη,..." (zit. nach der Ausgabe von BURNET 1962). Belege für die vorplatonische Bedeutung von ὁνομαζόμενος bzw. ὁνομαζόμενος finden sich bei BENVENISTE (1951). Sehr bekannt ist die spätere Definition von ARISTOXENUS, Rhythmica 1, fr. 1: "ἡ πρός τῶν τάξεως (zit. nach der Ausgabe von WESTPHAL 1965, Bd.II:75). Dem griechischen ὁνομαζόμενος entspricht im Lateinischen der Terminus *numerus*. Vgl. hierzu CICERO Orator 20,67: "Quid quid est enim quod sub aurium mensuram aliquam cadit... numerus vocatur, qui Graece ὁνομαζόμενος dicitur." (zit. nach der Ausgabe von HUBBEL 1962). LOTZ (1974:964) setzt dagegen irrtümlicherweise *numerus* mit μέτρον gleich.
- 4 Dies ist z.B. bei der Geige und der Singstimme möglich. Auch CICERO berücksichtigt dieses Phänomen anscheinend nicht. Vgl. CICERO, De oratore 3,186: "Numerus in continuatione nullus est; distinctio et aequalium et saepe variorum intervallo rum percussio numerum conficit; quem in cadentibus guttis, quod intervallis distinguuntur, notare possumus, in omni praecipitante non possumus." (zit. nach der Ausgabe von WILKINS 1963). Wie ich bereits angedeutet habe, muß man zwischen verschiedenen Arten der Rhythmusbildung differenzieren. Man kann z.B. unterscheiden: 1. Frequenz der Signale konstant, Zeitintervalle gleich Null, Amplitude variabel; 2. Frequenz variabel, Zeitintervalle gleich Null, Amplitude konstant; 3. Frequenz konstant, Zeitintervalle größer Null und eventuell variabel, Amplitude konstant. CHATMAN (1965) scheint sich vor allem auf den Fall 3 zu beziehen. Möglicherweise meint CHATMAN jedoch auch, daß Rhythmus nur dann wahrgenommen werden kann, wenn die Ereignisfolge mindestens aus zwei durch den Hörer differenzierbaren Signalen besteht und zwischen diesen unterschiedlichen Signalen ein Zeitintervall größer Null besteht. In diesem Fall muß allerdings nach BOLTON (1894) der Zeitabstand zwischen je zwei unterschiedlichen Signalen größer oder gleich 115 msec bzw. kleiner oder gleich 1580 msec betragen. Dies bedeutet, daß die periodische Wiederholung weder zu schnell erfolgen darf (dann tritt eine sog. Maskierung ein), noch zu langsam (dann wird phänomenal keine Gestalt gebildet). Aus dieser Sicht gehört eine Bestimmung des Zeitintervalls mit zu den notwendigen Konstituenten des Rhythmus.

- 5 Vgl. auch rhythmisches Arbeiten wie Mähen oder Hammerschlag des Schmiedes. Auch bei 'nonsense'-Wörtern scheint in der Kindersprache eine Tendenz zur Rhythmisierung vorzuliegen (vgl. OHNESORG 1966:56f). FUCS (1968:139) hat zur Überprüfung der These "Kein Mensch kann etwas gänzlich Ungeordnetes machen" Versuchspersonen aufgefordert, 1000 Zahlen (1-9) so unregelmäßig wie möglich nacheinander zu nennen. Es ergab sich eine signifikante Abweichung von einer Zufallsfolge.
- 6 Bei den Intervallen hat CHATMAN (1965:21) eine Spanne von 14,5 % festgestellt, bei den "events" gibt FRAISSE (1956:92) eine Spanne von 5 % an. Der Minimalwert der Intervalle ist nach CHATMAN (1965:19) 1/10 Sekunde, der Optimalwert nach KAYSER (1967:247) 2/3 Sekunde; weitere z.T. erheblich differierende Angaben von Minimal-, Maximal- und Optimalwerten beim Rhythmus finden sich bei CHATMAN (1965:21), FRAISSE (1956, passim) und WALLIN (1901, passim).
- 7 Vgl. auch JACOB (1918:181) und die Feststellung von CHATMAN (1965:22): "What is important is the impression of proportion or equivalence, not mathematically exact proportion or equivalence itself."
- 8 Neben kognitiven Faktoren sind bei der Perzeption rhythmischer Strukturen vermutlich auch affektive und vor allem motorische Faktoren involviert (vgl. FRAISSE 1956:5). Für STETSON (1903) ist das motorische Element sogar primär. Die gleiche Ansicht vertritt auch der Phonetiker ABERCROMBIE in bezug auf den Sprach- bzw. Versrhythmus: "The rhythm of speech, therefore, is primarily muscular rhythm, rhythm of bodily movement, rather than a rhythm of sound. This is why verse can be immediately recognized and felt as verse in silent reading, which otherwise would not be easy to explain." (1964:8)
- 9 Dies gilt übrigens, wie ich bereits im Zusammenhang mit der Theorieabhängigkeit von Beobachtungsdaten festgestellt habe, für jegliche Art der Wahrnehmung: "Our perceptions are a function of the nature of the stimuli, but also the 'assumptions' with which we apprehend them. This assumption itself depends on our previous experience, on the context of the perception, and on our personality." (FRAISSE 1963:145)
- 10 Vgl. auch FRY (1970:368): "Any spoken message embodies a series of different stresses just as it contains different tones. Every syllable in the message will bear some degree of stress and the succession of these stresses makes up a rhythmic pattern."
- 11 Vgl. COHEN (1966:90): "rythme = périodicité perçue"
- 12 Diese Ansicht wird schon in der Antike vertreten. Vgl. AUGUSTINUS: De musica V, 1: "Interesse igitur animadverterunt [docti veteres] inter rhythmum et metrum aliquid, ut omne metrum rhythmus, non etiam, omnis rhythmus metrum sit." (zit. nach der Ausgabe von FINAERT/THONNARD (1947:294). Häufig wird in der griechisch-lateinischen Metrik jedoch auch unter *metrum*(μέτρον) die kleinste rhythmische Einheit von ein oder zwei Versfüßen verstanden. Auf diese Bedeutung des Terminus Metrum wird im folgenden nicht weiter eingegangen.

- 13 Das Verhältnis von Metrum und Rhythmus und vor allem die Frage, inwieweit das Metrum als Abstraktion zu betrachten sei, hat zu heftigen Kontroversen geführt. Vgl. z.B. WIMSATT/BEARDSLEY (1959; 1964); PACE (1959); SCHWARTZ/WIMSATT/BEARDSLEY (1962).
- 14 Vgl. auch KAYSER (1967:97): "Unter M e t r u m versteht man das Schema eines Gedichts, das unabhängig von der sprachlichen Erfüllung existiert."
- 15 Vgl. dagegen LOTZ (1974:969), der Metrum und numerische Regularität gleichsetzt: "...numerical regularity, or meter, is the *differentia specifica* of verse."
- 16 Vgl. CHATMAN (1964:149): "Meter might be defined as a systematic literary convention whereby certain aspects of phonology are organized für aesthetic purposes." Die Konventionalität des Metrums und anderer metrischer Phänomene, wie z.B. Verszeile, Reim usw., wird vor allem von LEVIN (1971) betont.
- 17 TRUBETZKOY (1967:67) spricht von "merkmaltragend" und "merkmallos". Die Termini 'markiert' (*marked*) und 'unmarkiert' (*unmarked*) sind vor allem in der Markiertheitstheorie der generativen Phonologie üblich. Dort wird der Terminus 'markiert' zur Bezeichnung von Klassen von normalerweise auftretenden (frequenten, natürlichen) Beobachtungsdaten, 'unmarkiert' zur Bezeichnung von Klassen von ausnahmsweise auftretenden (weniger frequenten, weniger natürlichen) Beobachtungsdaten verwendet. Vgl. RÖŽIČKA (1970), BECHERT (1971).
- 18 Vgl. auch BERNHART (1974:115), der die Definition der *metricality* in der generativen Metrik kritisiert und feststellt: "For what is really decisive for metricality is the overall aesthetic effect of a line on the listener's mind." Allerdings ist m.E. der Ansatz von BERNHART ebenfalls unbefriedigend, da BERNHART keine deutliche Trennung zwischen perzeptorischer Ebene und Textbeschreibungsebene vornimmt.
- 19 Die von KLOPSTOCK geschaffenen 'Freien Rhythmen' werden von KAYSER (1967:96) folgendermaßen charakterisiert: "Der 'Freie Rhythmus' ist durch das Fehlen jeglicher metrischer Vorschrift gekennzeichnet: es gibt weder Reim, noch feste Strophen, noch feste Zeilen oder eine festliegende Füllung der Senkungen. Was von der Prosa unterscheidet, ist lediglich die Wiederkehr der Hebungen in annähernd gleichen Abständen." Zu weiteren Übergangsformen vgl. z.B. SUTTERHEIM (1961), FOWLER (1966a) und ERN (1968).
- 20 Vgl. LEVIN (1971:197): "The line is a purely typographic device of poetry." Vgl. auch die anhand von Textbeispielen begründete Ansicht von WODE (1970:391): "Gedichte haben auf der linguistischen Ebene keine Zeilen." Zum Verhältnis von sprachlicher Struktur und Vers(zeile) vgl. THOMPSON (1961).

- 21 Dieses Urteil muß natürlich relativ gesehen werden, da die (metrische) Kompetenz des Lesers bzw. Hörers eine Rolle spielt. Zu den sprachlichen Konstituenten des *vers libre* vgl. CZERNY (1961) und KLOEPFER (1971).
- 22 Vgl. jedoch z.B. GUÉRON (1974), die vom Text ausgeht und den Vers sprachlich mit Hilfe des Reims bestimmt.
- 23 Die Analyse von Verskonstituenten wie Reim oder Alliteration, oder auch die Analyse von Strophen wird häufig nicht als Aufgabe der Metrik, sondern als Aufgabe einer Verswissenschaft bzw. Verstheorie angesehen (vgl. WIMSATT 1972:XIX). Die Aufgabe der Metrik ist von diesem Standpunkt aus dagegen primär in der Analyse des Metrums zu sehen. Vgl. auch IHWE (1971/72, Bd. III), der Arbeiten wie LEVY (1965) oder LUELSDORFF (1968) zur Verstheorie, Arbeiten wie JAKOBSON/LOTZ (1941) oder BEAVER (1968a) dagegen zur Metrik rechnet.
- 24 Treffend beschreibt der niederländische Literaturwissenschaftler und Schriftsteller A. VERWEY das Verhältnis von Metrum und Rhythmus mit folgendem Bild: "Die Metra sind eigentlich entsetzliche Tote. Keiner sieht sie, wir sehen nur Rhythmus. Aber sie sind die Mathematik hinter den Rhythmen, die unerbitterliche, unkörperliche Gesetzmäßigkeit." (1934:62) Vgl. auch die Auffassung von WINTERS (1962:83): "Meter is the arithmetical norm, the purely theoretic structure of the line; rhythm is the controlled departure from that norm."
- 25 Vgl. auch SCHÄDLICH (1969:47f), der entsprechend der generativen Phonologie einen Unterschied macht zwischen phonologischer und phonetischer Repräsentationsebene sowie dem aktuellen (quasi-kontinuierlichen) Redesignal und das Verhältnis dieser drei Ebenen zu der JAKOBSONschen Trichotomie von "verse design" (Metrum), "verse instance" (rhythmische Struktur) und "delivery instance" (rhythmische Realisation) untersucht.
- 26 Vgl. z.B. ASMUTH (1976:221): "Unter Metrum (Versmaß) versteht man dabei das wiederholt vorkommende, also abstrahierbare, meist historisch vorgegebene Schema der Silben-, in längeren Versen auch der Kolonverteilung, unter Rhythmus seine individuelle Gestaltung oder auch deren Wiedergabe durch den Interpreten."
- 27 Vgl. die Charakterisierung von Metrum und Rhythmus bei GÁLDI (1967:692f): "Le mètre est impersonnel, le rythme, au contraire, personnel, subjectif, quelque chose de profondément vécu." Zur Abweichung vom Regelmäßigen als allgemeines Kunstprinzip vgl. SEASHORE/METFESSEL (1925:538): "In music and speech, pure tone, true pitch, exact intonation, perfect harmony, rigid rhythm, even touch, precise time play a relatively small role. They are mainly points of orientation for art and nature... This deviation from the exact is, on the whole, the medium for the creation of the beautiful - for the conveying of emotion. That is the secret of the plasticity of art. The exact is cold, restricted, and unemotional; and however beautiful, in itself soon palls upon us."

- 28 Vgl. auch HRUSHOVSKI (1964:179): "Poems with the same metrical scheme may have entirely different rhythms and can be rhythmically less alike than poems with different schemes, for example those written by one author."
- 29 Die psychische Leistung bei der Perzeption von metrischen Texten ist bereits von ALDEN (1914) betont worden. Vor allem beim "silent reading" ist nach ALDEN der Rhythmus weitgehend "mentally perceived".
- 30 Für BIERWISCH (1965:55) sind deshalb Strukturen wie Vers, Reim, Alliteration "parasitäre Strukturen, die nur auf der Grundlage linguistischer Primärstrukturen möglich sind."

6. Die Interdependenz von metrischen und sprachlichen Strukturen
- 1 Eine grundlegende Darstellung der Erklärungsproblematik findet sich bei STEGMÜLLER (1969). Vgl. auch RESCHER (1970); LENK (1972); GIESEN/SCHMID (1975). Zum Verhältnis von Erklärung zu (hermeneutischem) Verstehen vgl. VON WRIGHT (1974)
 - 2 Entsprechend den Adäquatheitsbedingungen liegt z.B. im folgenden Fall *k e i n e* Erklärung vor: BARTSCH/VENNEMANN (1973:35) meinen, daß "Allgemeinaussagen" wie "Wenn x ein Artikel und y ein Nomen ist, dann ist xy ein Nominalausdruck" und "Wenn x ein Nominalausdruck und y ein Verb ist, dann ist xy ein Satz" es erlauben, "den Satzcharakter einer großen Zahl von Wortfolgen nachzuweisen." Bei den Allgemeinaussagen handelt es sich jedoch nicht um empirische Gesetze, sondern um Definitionen. Ein solcher "Nachweis" ist daher auch keine erfahrungswissenschaftliche Erklärung wie offenbar zuweilen angenommen wird (vgl. WANG 1972). Zur Erklärung in der Linguistik vgl. auch WUNDERLICH (1974:92ff) und ANDRESEN (1974), die den Erklärungsgehalt linguistischer Theorien sehr gering einschätzt.
 - 3 Ich beziehe mich auf die Unterscheidung zwischen "sistema", "norma" und "habla" bei COSERIU (1952).
 - 4 CHATMAN (1965:96): "Meter itself is a system, parallel to and actualized by, but not confused with, the linguistic system."
 - 5 Vgl. auch die (Lieder)dichtung einiger Indianersprachen. So kommen z.B. in einigen Nootka-Liedern Phoneme vor, die die Normalsprache nicht aufweist (vgl. HOCKETT 1960:557). In Hopi-Liedern kommen akzentuelle Merkmale vor, die in der Normalsprache nicht verwendet werden (vgl. VOEGELIN/EULER 1957). Weitere Hinweise finden sich bei HERZOG (1946).
 - 6 Vgl. z.B. *pōpulus* (Volk) : *pōpulus* (Pappel). Vgl. auch das Persische, das Arabische und das klassische Türkische.
 - 7 Vgl. Wortpaare wie dt. Übersetzen : Übersetzen; engl. *to permit* : *permit*
 - 8 Zum Binarismus in der Linguistik vgl. HENRICI (1975). Nach JAKOBSON/HALLE (1970:444) scheint die Binarität ein Teil der psychischen Grundkonzeption der Menschen zu sein.
 - 9 Vgl. COLLINDER (1951:40); JONES (1962:67). OKSAAR (1971) nimmt für das Estnische sowohl bei den Vokalen als auch bei den Konsonanten drei Quantitäten an.
 - 10 Weitere Beispiele finden sich u.a. bei KAYSER (1967:91f), STANKIEWICZ (1964).

- 11 Vgl. vor allem die Arbeit von KURYŁOWICZ (1975), der das Verhältnis von Metrik und Sprachgeschichte anhand von mehr als zehn metrischen Systemen analysiert. STANKIEWICZ (1964) vergleicht das russische und das polnische metrische System und betont den Einfluß von kulturellen Faktoren. Wichtig sind in diesem Kontext auch die Arbeiten von KIPARSKY (1968, 1972), der anhand des Kalevala und des Rigveda u.a. die Bedeutung der diachronen Untersuchung metrischer Texte für die Überprüfung von Hypothesen der generativen Phonologie aufzeigt.
- 12 Ich lasse in diesem Zusammenhang die sehr umstrittene Alternationstheorie unberücksichtigt. Vgl. hierzu HUBER (1963), SCHMIDT (1968) und VERRIER (1931-1932), der wohl als einziger französischer Versteoretiker die Alternationstheorie vertritt. Die Funktion des Reims als metrische Konstituente dürfte auch der Grund dafür sein, daß der Reim im Französischen eine wichtigere Rolle spielt als im Deutschen. So blieb auch der Versuch von GUÉRIN (Song des Crépuscules, 1895), den Reim durch Assonanz zu ersetzen, ohne dauernden Erfolg (vgl. KAYSER 1967:97).
- 13 Es gibt nur wenige quantitative Oppositionen im Französischen wie z.B. *bête* - *bette* oder *pâte* - *patte*, die jedoch von vielen Franzosen nicht mehr beachtet werden (vgl. LÉON 1966:19f).
- 14 Auch im englischen metrischen System hat sich der *quantitative verse* trotz wiederholter Versuche nicht endgültig durchsetzen können. Vgl. hierzu HASCALL (1974:65).

7. Die Silbe als metrische Konstituente

- 1 Es ist oft problematisch, zwischen konstitutiven und nicht-konstitutiven Elementen zu unterscheiden. Ich betrachte solche Elemente als konstitutiv, die zur Konstitution der Metrizität eines Textes beitragen. So ist z.B. im Altgermanischen die Alliteration ein wichtiges Mittel, um die Kohärenz der Langzeile zu gewährleisten, während im Französischen die Alliteration im Gegensatz zum Reim nicht zur Konstitution des Iso-syllabismus beiträgt. Eine andere Sicht des Verhältnisses von konstitutiven und nicht-konstitutiven Elementen findet sich bei LOTZ (1964:139f).
- 2 Bekannt sind die häufigen Parallelismen des Alten Testaments, besonders der Psalmen (vgl. auch NEWMAN/POPPER 1918, ALBRICHT 1950, CROSS 1950). Ein Blick auf andere Literaturen zeigt die Universalität des Parallelismus vor allem in epischer Sprache (vgl. z.B. STEINITZ 1934, AUSTERLITZ 1958, POPPE 1958, GONDA 1959).
- 3 Dies gilt auch für eine Vielzahl von (sonstigen) linguistischen Arbeiten (vgl. PULGRAM 1970:11).
- 4 Vgl. die umfangreichen Bibliographien z.B. bei TILLMANN (1964), GRANDA GUTIERREZ (1966), PIKE (1967), PULGRAM (1970), HÁLA (1973).
- 5 Gerade in letzter Zeit ist im Rahmen der generativen Phonologie auf die Wichtigkeit des Silbenbegriffs hingewiesen worden. So hat z.B. VENNEMANN GENANNT NIERFELD (1972) gezeigt, daß sich viele phonologische Prozesse erst mit Hilfe des Silbenbegriffs erklären lassen. Vgl. auch LÜDTKEs (1970) kybernetisches Sprachmodell, in dem die Silbe im Gegensatz zum Phonem, das als "fiktive sprachliche Einheit" (S. 39) eliminiert wird, eine zentrale Rolle spielt. Auch die Frage, ob die Silbe als sprachliches Universale anzusehen ist, ist kontrovers (vgl. z.B. KOHLER 1966: 208, CHERRY 1967:197, FUDGE 1969:253).
- 6 LEHISTE (1970:155), die sich auf FROMKIN (1968) stützt, vermutet ein solches Korrelat: "One unit of neural organization seems to be the syllable (FROMKIN 1968). There is strong evidence for this from a number of areas, production as well as perception." Weitere Hinweise zur psychischen Realität der Silbe finden sich z.B. bei LEUNINGER/MILLER/MÜLLER (1972:18ff). Auch die intuitive Silbenvorstellung des Menschen, die Verwendung von Silbenschriften wie Linear B und Ergebnisse der Aphasieforschung (vgl. z.B. JAKOBSON 1969:87) sprechen für die psychische Realität der Silbe.
- 7 Vgl. auch PILCH (1968:18): "Akustisch weist jede Silbe ein Intensitätsmaximum auf." Zur Silbenintensität vgl. auch FLETCHER (1965:84).

- 8 Vgl. auch die Arbeiten von KURYŁOWICZ (1948), FISCHER-JØRGENSEN (1952), HAUGEN (1956) und die Feststellung von KLOSTER-JENSEN (1963:34): "Es ist dringend notwendig, daß man sich klar macht, was für eine Hybride die Silbe ist. Sie hat einen phonetischen Kern und phonemisch und distributionell bestimmte Grenzen."
- 9 Vgl. z.B. für das Russische BALDWIN (1969) und für das Englische O'CONNOR/TRIM (1953) und für das Französische MOREAU (1966).
- 10 Auch bei der Bestimmung von positionslangen Silben im Lateinischen und Griechischen spielt die Frage der Silbengrenzen keine Rolle, da die Positionslänge lediglich von der Anzahl der dem Silbengipfel folgenden Konsonanten abhängt, nicht jedoch von der Silbengrenze. Vgl. auch BERSCHIN (1970) und vor allem die zahlreichen Arbeiten von OTT (z.B. OTT 1973), die zeigen, daß für die automatische Analyse des lateinischen daktylischen Hexameters keine exakte Bestimmung der Silbengrenzen nötig ist. Ausführlich wird das Problem der Silbendelimitierung im Lateinischen bei ZIRIN (1970) diskutiert.
- 11 Vgl. auch LOTZ (1974:972f) und HALLE/KEYSER (1972:234) Anmerkung 1: "We use the term 'syllabic' here as the equivalent of 'sequence of speech sounds consisting of one syllabic sound ('vowel') preceded and followed by any number of consecutive nonsyllabic sounds ('consonants')." In particular, we do not take a position on the vexing question of whether or not utterances can be unambiguously segmented into syllables."
- 12 Bereits in der Antike findet sich in bezug auf die Silbe eine 'funktionsorientierte' und eine 'lautorientierte' Auffassung. Vgl. z.B. CHOROBOSCUS Kommentar zum Enchiridion des HEPHAESTION (zit. nach der Ausgabe von CONSRUCH 1906:180):
 'Ιστέον δὲ ὅτι ἄλλως λαμβάνουσι τοὺς χρόνους οἱ μετρικοί, ἤγουν οἱ γραμματικοί, καὶ ἄλλως οἱ ῥυθμικοί. οἱ γραμματικοὶ ἐκείνων μακρὸν χρόνον ἐπίστανται τὸν ἔχοντα δύο χρόνους, καὶ οὐ κατὰγονται εἰς μετρὸν τι· οἱ δὲ ῥυθμικοὶ λέγουσι τὸδε εἶναι μακρότερον τοῦδε, φάσκοντες τὴν μὲν τῶν συλλαβῶν εἶναι δύο ἡμίσεος χρόνων, τὴν δὲ τριῶν, τὴν δὲ πλειόνων· οἷον τὴν $\overline{\omega}$ ς οἱ γραμματικοὶ λέγουσι δύο χρόνων εἶναι, οἱ δὲ ῥυθμικοὶ δύο ἡμίσεος, δύο μὲν τοῦ $\overline{\omega}$ μακροῦ, ἡμιχρονίου δὲ τοῦ σ .
 "Man muß wissen, daß die Metriker oder Grammatiker Zeiteinheiten anders auffassen als die Rhythmiker. Die Grammatiker meinen, daß die lange Zeiteinheit zwei Zeiteinheiten enthalte, und sie führen sie nicht auf etwas größeres zurück. Aber die Rhythmiker sagen, daß eine lange Zeiteinheit länger sei als eine andere lange Zeiteinheit, und sie behaupten, daß einige der Silben zwei und eine halbe Einheit lang seien, andere drei und andere noch mehr. So sagen z.B. die Grammatiker, daß $\overline{\omega}$ ς zwei Einheiten lang sei, während die Rhythmiker sagen, es sei zwei und eine halbe Einheit lang, zwei Einheiten für das lange $\overline{\omega}$ und eine halbe für das σ ." (Übersetzung R.G.)

8. Der Reim
- 1 Vgl. auch ABERNATHY (1967), der die eingeschränkte Gültigkeit des Reims als Äquivalenzrelation u.a. mit Beispielen aus dem Deutschen und Russischen belegt.
- 2 Vgl. ABERNATHY (1967), der die Elemente von $\bar{R}$ *non-rhymes* bzw., wenn die Relation $\bar{R}$ konventionalisiert ist, *antirhymes* nennt.
- 3 Die folgenden Definitionen entsprechen z.T. den Definitionen bei LUELSDORFF (1968). LUELSDORFF berücksichtigt jedoch in seinen Definitionen lediglich den sog. reinen Reim. Ein weiterer Formalisierungsvorschlag findet sich bei VAN DIJK (1972:217).
- 4 JAKOBSON (1964:368) betrachtet deshalb den Reim als eine Sonderform des weit allgemeineren poetischen Phänomens des Parallelismus.
- 5 Der Terminus 'Konnotation' wird je nach linguistischer Schule in verschiedener Bedeutung gebraucht. Außerdem wird dieser Terminus auch in der Psychologie und in der Logik verwendet (vgl. GARY-PRIEUR 1971; MOLINO 1971). Es sollen im folgenden unter 'Denotationen' primäre, objektive, einer geschlossenen Liste angehörende Bedeutungen, unter 'Konnotationen' dagegen sekundäre, individuelle, einer offenen Liste angehörende Bedeutungen verstanden werden (vgl. GARY-PRIEUR 1971:99).
- 6 TURČANY (1964) spricht in diesem Kontext auch von einer "ver-tikalen Metapher".
- 7 zit. ELWERT (1966:98).
- 8 zit. ELWERT (1966:95).
- 9 Der Reim kann somit zu dem für die meiste (heutige) Poesie typischen Streben nach Ambiguität beitragen. Aufgrund ihrer andersartigen Auffassung von der Funktion der Dichtung haben sich die französische Klassik und vor allem die französische Aufklärung gegen den Reimzwang gewehrt. Vgl. BOILEAU, Art poétique I, 30: "La Rime est une esclave, et ne doit qu'obéir." Vgl. die Kritik am Reim z.B. bei VOLTAIRE: Sur Oedipe, Lettre 5 (BEUCHOT 1830, Bd. II:45). Weitere Belege finden sich bei VAN TIEGHEM (1965:94).
- 10 Vgl. MALMBERG (1967:164): "Rhyme and alliteration imply identity between phonemes not between allophones." Damit fordert MALMBERG phonologische Identität für den Reim überhaupt und berücksichtigt dabei nicht den unreinen Reim. Vgl. auch KRÄMER (1971:13).
- 11 Vgl. z.B. italienisch *parlo* : *parlò*; deutsch Übersetzen : Übersetzen; spanisch *límite* : *límite* : *limité*.
- 12 Es besteht auch die Möglichkeit, die letzte Silbe von "Vorstellung" zu akzentuieren und somit der Betonung von "Schwung" anzugleichen. In der Regel wird jedoch eine solche gegen die normale Akzentuierung verstoßende Diktion als unangemessen abgelehnt.

- 13 Vgl. auch Reime wie *croce* : *conduce* und *dire* : *vedere*, die bis zum *dolce stil nuovo* im Italienischen üblich waren. Auch diese Reime beruhten auf sizilianischem Vorbild, wo *conduci* : *cruci* und *diri* : *vidiri* lautgerecht reimten (ELWERT 1968:80f).
- 14 Eine Überprüfung der Berechnungen von KRÄMER ergab sogar nur 11,86 % nicht-identische Reime. Denkt man an die Entwicklung des slowakischen Reims in der Zeit von 1840 - 1960, die durch eine Abnutzung und Banalisierung des offenen Reims zu Gunsten des geschlossenen Reims gekennzeichnet ist (vgl. STUKOVSKÝ/ALTMANN 1966), liegt die Hypothese nahe, daß aufgrund einer allmählichen Abnutzung des identischen Reims der nicht-identische Reim vor allem bei jüngeren Autoren zu finden ist. Diese Hypothese läßt sich anhand der acht von KRÄMER untersuchten Autoren nicht bestätigen. Denn berücksichtigt man den Stichprobenfehler, ergeben sich auf der Basis der Daten von KRÄMER (1971:15f) folgende Autorengruppen: 1) TRAKL (31,33 %) 2) MÖRIKE, HÖLDERLIN, BUSCH (16,46 %) 3) NIETZSCHE (13,64 %) 4) MORGENSTERN (6,27 %) 5) RILKE, RINGELNATZ (3,68 %). Allerdings sind vermutlich die Variablen 'Genrestil' und 'Individualstil' von erheblicher Bedeutung, so daß eine Erhöhung der Anzahl der untersuchten Autoren - STUKOVSKÝ/ALTMANN haben insgesamt 46 Autoren untersucht - zu einer Bestätigung der Hypothese führen könnte.
- 15 Es gibt mittlerweile eine Reihe von Vorschlägen zu distinktiven Merkmalen. Am bekanntesten sind die von JAKOBSON (vgl. JAKOBSON/FANT/HALLE 1951) entwickelten Merkmale (modifiziert in: JAKOBSON/HALLE 1970). Neuerdings werden auch häufig die von CHOMSKY/HALLE (1968) vorgeschlagenen Merkmale verwendet.
- 16 Vgl. auch die veränderte phonologische Struktur der poetischen Sprache bei einigen Indianersprachen, wo sich Alltagssprache und poetische Sprache in einigen distinktiven Merkmalen unterscheiden können (vgl. Kap. 8).
- 17 Im Englischen z.B. gibt es wenigstens vier oder fünf Arten von [t]-Lauten, mit deutlichen akustischen Unterschieden. Diese Unterschiede werden von einem muttersprachlichen Engländer nicht bemerkt, während ein Inder einige Unterschiede hört, da diese in seiner Muttersprache phonologisch relevant sind (vgl. FRY 1971:34).
- 18 Zur extensionalen Bestimmung vgl. z.B. UNGEHEUER (1965) oder HANSON (1967).
- 19 Vgl. auch KRÄMER (1971:26), der vermutet, daß eine konsonantische Differenz weniger deutlich perzipiert wird und der entsprechende Reim deswegen als phonemisch identisch beurteilt werden kann.
- 20 Dies ergab ein informeller Perzeptionstest mit 10 Vpn.
- 21 LÜDTKE (1953) nimmt für das Portugiesische 13 Vokalphoneme an. Zur Anzahl der Diphthonge und Triphthonge vgl. LINDSTRAND (1943).
- 22 Wenn wir für die Dichtung die Existenz eines /a/ *postérieur* und des Nasals /ɔ̃/ annehmen, erhalten wir 16 Vokalphoneme (vgl. KLEIN 1966:44).

- 23 Zum Platz des Akzents im Deutschen und Spanischen vgl. DELATTRE (1965:29), zum Italienischen LICHEM (1969:122).

9. Metrik und Sprachtypologie

- 1 Vgl. z.B. LEVÝ (1961), JAKOBSON (1966a), BECK (1966), VALESIO (1966), KURYŁOWICZ (1975).
- 2 Vgl. z.B. die Bemerkungen von STANKIEWICZ (1964:79) zum Englischen, Russischen und Deutschen.

10. Die Funktion metrischer Strukturen

- 1 Vor allem in der Antike wurde die Auffassung vertreten, daß metrische Strukturen eine schmückende Funktion (*ornatus*, ὀρεός) haben. Zur rhetorischen (persuasiven) Funktion metrischer Strukturen vgl. SIEVEKE (1972). Eine Reihe von Hinweisen speziell zur Funktion des Metrums finden sich bei CHATMAN (1965:184ff). IHWE (1975:395f) sieht die Funktion metrischer Texte in der Einführung und Konstruktion spezifischer Typen von Diskurswelten.
- 2 Bei einer umfassenderen Analyse der Funktion metrischer Strukturen müßten auch die verschiedenen Sprachfunktionen berücksichtigt werden. BÜHLER (1965:28) hat in seinem Organonmodell der Sprache drei Funktionen zugewiesen: die "Ausdrucks"-, "Appell"- und "Darstellungsfunktion" ("symptomatische, appellative, symbolische Funktion"). Vgl. auch BÜHLER (1918:1). JAKOBSON (1964:352) unterscheidet sechs verschiedene Sprachfunktionen. In formalisierter Darstellungsweise finden sich diese sechs Funktionen auch bei SEBEEK (1963:53), der auch auf den Unterschied zum BÜHLERSchen Organonmodell eingeht. Auch bei KAINZ (1967:172ff) und bei LEONT'EV (1971:33ff) findet sich eine über das Organonmodell hinausgehende Darstellung der Sprachfunktionen. Zum Problem der Sprachfunktionen im Prager Strukturalismus (Funktionalismus) vgl. CERVENKA (1973).
- 3 In diesem Kontext müßten auch die Ergebnisse der experimentellen Ästhetik (vgl. z.B. BELLYN 1975) und der Biophonetik (TROJAN 1975) berücksichtigt werden. Die Untersuchung der Wirkung metrischer Strukturen ist auch für eine Theorie der (literarischen) Übersetzung von großer Bedeutung. Denn nach REISS (1971:39) ist das oberste Gebot bei der Übersetzung künstlerischer Texte, die gleiche ästhetische Wirkung zu erzielen (vgl. auch ALBRECHT 1973:13).
- 4 Geht man vom Organonmodell BÜHLERS aus, könnte die Funktion der lautlichen Form auch als Kundgabe (Ausdruck) und Appell interpretiert werden. Eine "Kundgabe"- und "Appellphonologie" (vgl. VON LAZICZIUS 1935, TRUBETZKOY 1967:19) ist allerdings bisher erst in Ansätzen vorhanden (vgl. MALMBERG 1971:9). Vgl. auch STANLEY (1971:95), der von "phonopragmatics" und "phonosemantics" spricht. Die im folgenden untersuchten Phänomene können zum großen Teil auch zum Objektbereich der sog. Paralinguistik gerechnet werden.
- 5 BIERWISCH (1965:63) schlägt vor, Struktureigenschaften, die im Aufnahmebereich des Kurzzeitgedächtnisses liegen, Mikrostrukturen, solche die außerhalb des Kurzzeitgedächtnisses liegen, Makrostrukturen zu nennen.
- 6 Die meisten Arbeiten bleiben weitgehend spekulativ. Ein illustratives Beispiel hierfür ist die 'Analyse' der berühmten VERLAINE-Verse
Il pleure dans mon coeur ...
durch den bekannten französischen Metriker GRAMMONT (1923:366f).

Ein interessantes Pendant zu der klanglichen Expressivität von Lauten ist deren visuelle Symbolik, wie sie uns in dem Sonett "Voyelles" von RIMBAUD begegnet:

A noir, E blanc, I rouge, U vert, O bleu: voyelles
Je dirai quelque jour vos naissances latentes:

(zit. nach der Ausgabe von BERNARD 1960:110).

- 7 Neben den Arbeiten von FÓNAGY sind vor allem auch die Veröffentlichungen von DELBOUILLE (vgl. z.B. DELBOUILLE 1961, 1967) bekannt geworden. Zur Einschätzung der Arbeiten von DELBOUILLE und FÓNAGY vgl. MELHEM (1973). Vgl. auch GRAUBNER (1976:183), der die semantische Information z.B. von Lauten, Lautverbindungen oder Rhythmus mit dem Terminus "phonetische Konnotation" bezeichnet.
- 8 Vgl. z.B. die Hinweise bei HÖRMANN (1970) und PETERFALVI (1970).
- 9 So meint z.B. DE SAUSSURE (1968:100): Ce principe de l'arbitraire du signe n'est contesté par personne." (vgl. hierzu ENGLER 1962, 1964). Eine ähnliche Meinung vertritt LYONS (1969:54f): "Few people these days would maintain that the correlation of a particular word and a particular meaning is other than conventional. The long controversy between the 'naturalists' and the 'conventionalists' may be considered closed." Die Frage, ob zwischen der äußeren Gestalt des sprachlichen Zeichens und dem Inhalt ein natürlicher (φύσει) oder ein konventioneller (θέσει, νόμῳ) Zusammenhang besteht, war bereits in der Antike kontrovers (vgl. STEINTHAL 1961:41ff). Einen analogen Streit gab es in Indien (vgl. REGNAUD 1884:12ff). Weitere Belege für die bis zum heutigen Zeitpunkt fortgesetzte Diskussion finden sich z.B. bei FÓNAGY (1971) und KUTSCHERA (1975). Das deutlichste Beispiel für nicht-arbiträre Zeichen sind die Onomatopoetika, denen auch DE SAUSSURE einen Sonderstatus zubilligt.
- 10 Zit. nach der Ausgabe von MONDOR/JEAN-AUBRY (1965:364). Vgl. hierzu JAKOBSON (1969a:242).
- 11 Vgl. LÉVI-STRAUSS (1957:106f): "Si nous admettons donc, conformément au principe saussurien, que rien ne prédestine a priori certains groupes de sons à désigner certains objets, il n'en semble pas moins probable qu'une fois adoptés, ces groupes de sons affectent de nuances particulières le contenu sémantique qui leur est lié ... Quand nous envisageons le vocabulaire a posteriori, c'est à dire déjà constitué, les mots perdent beaucoup de leur arbitraire, car le sens que nous leur donnons n'est plus fonction seulement d'une convention." Zuweilen wird die Lautsymbolik auch mit Hilfe semiotischer Kategorien beschrieben. Expressive Phoneme z.B. werden als Ikone aufgefaßt. Vgl. z.B. MORRIS (1955:191): "An iconic sign, it will be recalled, is any sign which is similar in some respects to what it denotes. Iconicity is thus a matter of degree. It can obviously be a property of auditory and visual signs alike. Spoken language contains some sounds which

are clearly iconic ('onomatopoetic'). Some linguists have claimed this quality for certain vowels, and the poet certainly at times reproduces in the rhythm of his words movements of the objects which the poem signifies." Vgl. hierzu auch NÖTH (1976).

- 12 Als illustratives Beispiel für die auch in neueren Arbeiten zur Metrik noch des öfteren zu findende umfassende 'hermeneutische' Deutung der Funktion des Rhythmus sei SCHULTZ (1972:46) zitiert: "Der Rhythmus ist das Medium, das Vergangenheit und Zukunft verbindet. Keine rhythmische Härte darf im romantischen Lied den Strom der sehnsüchtigen 'Erinnerung' stören, keine Spannung zwischen Satz- und Versrhythmus darf auftreten. Der Hörer soll allmählich hineingenommen werden in die unendliche Geschichtsbewegung. Zugleich wird die Distanz zwischen Ich und Welt aufgehoben, wird das lyrische Ich mit der Natur eins."
- 13 Zit. nach der Ausgabe von CARLOS/ROBLES (1961, Bd. II:895).

11. Zur Theorie der mathematischen Metrik

- 1 Der Terminus 'mathematische Metrik' ist bisher noch nicht gebräuchlich. HERDAN (1969) z.B. verwendet den Terminus "mathematische Verstheorie"
- 2 Dies gilt nicht oder nur sehr bedingt z.B. für generative Ansätze, die Tiefenstrukturen postulieren oder für die (algebraische) Theorie der vagen Mengen. Einen Überblick über die für die Anwendung wichtige Theorie der vagen Mengen gibt KAUFMANN (1975).
- 3 Mathematische Metrik und mathematische Poetik unterscheiden sich vor allem dadurch, daß die mathematische Poetik primär am künstlerischen Verständnis von Texten interessiert ist (vgl. z.B. MARCUS 1973:282). Zum Verhältnis von mathematischer Texttheorie und Literaturwissenschaft vgl. FISCHER (1970b).
- 4 Vgl. z.B. die Arbeiten des englischen Statistikers HERDAN und des deutschen Physikers FUCKS. Weitere Hinweise zur Geschichte der mathematischen Metrik finden sich bei LEVY (1965). Insgesamt gesehen ist speziell im Bereich der quantitativen Metrik die Literaturlage kaum noch übersehbar (vgl. zur Literaturlage MARCUS 1973:211, LOTMAN 1973:178). Die quantitative Metrik dürfte auch der am meisten bearbeitete Bereich der quantitativen Textanalyse sein.
- 5 Vgl. auch ABRAHAM/BRAUNMÜLLER (1971:1): "Stil ist eine Funktion pragmatischer Variablen". Einige der folgenden Faktoren werden meist als Performanzfaktoren bezeichnet. Zur Unterscheidung von pragmatischen Faktoren und Performanzfaktoren vgl. z.B. WUNDERLICH (1970) oder ABRAHAM/BRAUNMÜLLER (1971). Zum Verhältnis von Stilistik und Psycholinguistik vgl. SLAMA-CAZACU (1967), zum Verhältnis Stilistik und Soziolinguistik vgl. BENES/VACHEK (1971). Die im folgenden genannten Faktoren könnten auch zur Textbeschreibung mit Hilfe einer Varietätengrammatik (vgl. KLEIN 1974a) benutzt werden.
- 6 Ein solches Vorgehen kann als Ex-post-facto-Experiment (Quasi-Experiment) interpretiert werden (vgl. z.B. MAYNTZ/HOLM/HÜBNER 1972:186ff, KERLINGER 1973:378ff).
- 7 Die Stileigenschaften kennzeichnen wiederum die "stilistische Kompetenz" (ABRAHAM 1971) oder auch die "poetische Kompetenz" (BIERWISCH 1965) des Textproduzenten.
- 8 Es sind bisher eine Vielzahl verschiedener Stildefinitionen vorgeschlagen worden. Z.B. wird 'Stil' als Wahl zwischen verschiedenen sprachlichen Möglichkeiten oder als Abweichung von einer bestimmten Norm oder auch als Addition bestimmter stilistischer Eigenschaften definiert (vgl. SANDERS 1973:13ff, ENKVIST 1973:14ff). In statistischen Untersuchungen wird vor allem die Konzeption von 'Stil als Abweichung

von einer Norm' vertreten. Hier stellt sich jedoch das Problem, wie diese Norm zu bestimmen ist. BLOCH (1953:40) z.B. sieht die "language as a whole", DOLEZEL (1965:289) den "durchschnittlichen Text", PIEPER (1975:76) den "typischen Text" als Vergleichsnorm an. Die von mir vorgeschlagene Definition impliziert sowohl eine Konzeption von 'Stil als Wahl' als auch von 'Stil als Abweichung'. Denn die pragmatisch bedingte Wahl führt zu Differenzen zwischen Textcharakteristiken, und Differenzen implizieren Abweichungen. Der Bezugsrahmen für die Abweichung ist allerdings im Verhältnis zu den genannten Definitionen relativ frei und wird vom jeweiligen Untersuchungsziel bestimmt (vgl. MÜLLER 1972:161f; ENKVIST 1973:24). Zum Begriff der Abweichung in der Stilistik vgl. auch GUEUNIER (1969).

- 9 $T(Q_i)$ und $T(Q_k)$ bezeichnen im Kontext Q_j bzw. Q_k produzierte Klassen von Texten.
- 10 Vgl. die Verwendungsweise des Terminus 'stationär' in der Theorie der stochastischen Prozesse. Die stationären bzw. nicht-stationären Stilcharakteristiken kennzeichnen wiederum die Stationarität bzw. Nicht-Stationarität des Kodiervorgangs.
- 11 Da Hypothesen in der Regel im Zusammenhang mit bestimmten konkreten Forschungsproblemen formuliert werden, wird von manchen Autoren auch als erster Forschungsschritt die Formulierung von Forschungsproblemen genannt. Forschungsprobleme haben im Gegensatz zu Hypothesen die Form von Interrogativsätzen und sind meist allgemeiner als die entsprechenden Hypothesen. Im Zusammenhang mit der quantitativen Textanalyse könnte es sich z.B. um Fragestellungen folgender Art handeln: "Sind die 'Nachtwachen des Bonaventura' Fr. Gottlob WETZEL zuzuschreiben?", "Hat sich der Stil GOETHEs im Laufe seines Lebens verändert?", "Ist es möglich, mit Hilfe numerischer Kriterien zwischen metrischen und nicht-metrischen Texten zu differenzieren?", "Besteht ein Unterschied zwischen den Hexametern Vergilischer und Nicht-Vergilischer Dichter?" usw.
- 12 Kriterien zur Differenzierung zwischen verschiedenen Typen von Hypothesen finden sich z.B. bei BUNGE (1967 I:238ff u. 256) und LEINFELLNER (1967:104ff).
- 13 Es werden u.a. folgende Arten von Validität unterschieden: *predictive validity*, *concurrent validity*, *content validity* und *construct validity*. Vgl. hierzu z.B. GHISELLI (1964:336ff).
- 14 Es gibt eine Reihe verschiedener Verfahren, um die Validität und die Reliabilität eines empirischen Forschungsinstruments (Indikatoren, operationelle Definitionen, Meßinstrumente) zu bestimmen. Vgl. hierzu z.B. FRIEDRICH/HENNIG (1975).
- 15 Vgl. HRABÁK (1966:265): "In applying quantitative methods in the science of literature we run up against the difficulty that a good mathematician rarely understands poetry thoroughly and on the other hand a poetry expert is rarely a good mathematician."

- 16 Die folgenden Probleme sind Gegenstand der Meßtheorie (Skalierungstheorie). Vgl. hierzu z.B. TORGERSON (1958), SUPPES/ZINNES (1963), PFANZAGL (1968), GUTJAHR (1971), SCHEUCH/ZEHNPFFENNIG (1974). In der Linguistik (Textwissenschaft) sind bisher erst vereinzelt Skalierungsprobleme diskutiert worden (vgl. z.B. NIKITPOULOS 1972). Dagegen haben sich Disziplinen wie Psychologie oder empirische Sozialforschung schon relativ früh mit der Skalierungsproblematik auseinandergesetzt.
- 17 Ich verzichte auf eine exakte algebraische Charakterisierung dieser vier Skalentypen. Vgl. hierzu z.B. KLIEMANN/MÜLLER (1972:172ff). Die vier Skalentypen lassen sich auch aufgrund ihrer Transformationseigenschaften charakterisieren. Die Nominalskala ist invariant gegen eindeutige Transformationen, die Ordinalskala gegen monotone Transformationen, die Intervallskala gegen positiv lineare Transformationen, die Verhältnisskala gegen Ähnlichkeitstransformationen (vgl. TORGERSON 1958). Die Anzahl der zulässigen Transformationen verringert sich somit mit zunehmendem Skalierungsniveau.
- 18 Sehr häufig wird der Meßbegriff auch im Sinne von Skalierung oder Metrisierung verwendet. STEVENS (1959) z.B. gebraucht den Meßbegriff sowohl im Sinne von metrischer als auch nichtmetrischer Skalierung. GUTJAHR (1971) dagegen verwendet den Begriff des Messens lediglich im Sinne von metrischer Skalierung. Zwischen dem weiten Meßbegriff von STEVENS und dem engen Meßbegriff von GUTJAHR steht der Meßbegriff von TORGERSON (1958), der bereits eine Skalierung auf Ordinalniveau als Messen bezeichnet.
- 19 Für die Güte einer Skalierung gelten somit auch die Kriterien der Validität und Reliabilität.
- 20 Die Termini 'qualitativ' und 'quantitativ' werden nicht einheitlich verwendet. Z.B. werden die komparativen Begriffe von manchen Autoren auch zu den qualitativen Begriffen gezählt (vgl. KUTSCHERA 1972:24).
- 21 Ich gehe davon aus, daß a priori nichts gegen die Möglichkeit der Quantifizierung irgendeines empirischen Forschungsgegenstandes spricht. Zum Verhältnis von Qualität und Quantität vgl. z.B. ESSLER (1971:64ff). Hinweise auf weitere Vorteile der Quantifizierung (Metrisierung) finden sich bei GUTJAHR (1971:30f).

12. Quantitative Analyse metrischer Texte

- 1 LOTZ (1956:4) stellt irrtümlich in bezug auf den "Erlkönig" fest: "This poem is written in so-called Knittelvers."
- 2 Zur literaturwissenschaftlichen Einordnung und Interpretation des "Erlkönig" vgl. z.B. HIRSCHENAUER/WEBER (1968:159ff) oder BOYD (1961:170ff). Bei HIRSCHENAUER/WEBER finden sich auch Hinweise zu den zahlreichen Vertonungen des "Erlkönig". Es wäre möglicherweise aufschlußreich, die Vertonungen mit Hilfe mathematisch-statistischer Methoden zu untersuchen und die Ergebnisse mit den Resultaten aus der Erlkönig-Analyse zu vergleichen. Zur mathematischen Musikanalyse vgl. die Arbeiten der Gruppe um FUCKS (z.B. FUCKS 1963, FUCKS/LAUTER 1965a).
- 3 Die folgenden Bemerkungen zur 'Logik der Statistik' sind stark vereinfacht und mit erheblichen Problemen verbunden. Eine ausführliche Diskussion der angeschnittenen Fragenkomplexe findet sich bei STEGMÜLLER (1973a).
- 4 Dieses Verfahren haben auch ALTMANN/ŠTUKOVSKÝ (1965) bei der Analyse des malaisischen Pantun benutzt. Sie haben dabei festgestellt, daß die Zahl der langen Wörter im Verlauf des Verses zunimmt, während die Zahl der kurzen Wörter abnimmt.
- 5 Bei den Vpn handelt es sich um 12 Männer und 8 Frauen im Alter zwischen 22 und 62 Jahren mit vorwiegend universitärer Ausbildung. Wie eine zusätzliche Befragung ergab, war der Text des "Erlkönig" allen Probanden aus dem Schulunterricht bekannt.
- 6 Eine weitere Möglichkeit zur Untersuchung des Rezipientenverhaltens bietet die auf OSGOOD u.a. zurückgehende Technik des "semantischen Differentials" (vgl. z.B. OSGOOD 1952, OSGOOD/SUCI/TANNENBAUM 1957, SNIDER/OSGOOD 1969, ERTTEL 1969). Die einzelnen Strophen oder auch die einzelnen Personen des "Erlkönig" wären in einem dreidimensionalen Raum mit den Dimensionen (Koordinaten) Potenz (*potency*), Valenz (*evaluation*) und Erregung (*activity*) zu lokalisieren. Die Dimensionen beschreiben wiederum Bündel von bipolaren ordinalskalierten Adjektiven oder Substantiven wie z.B. Stärke, Härte... vs. Schwäche, Weichheit... (Potenz), anziehend, wohlklingend... vs. abstoßend, mißklingend... (Valenz) und schnell, erregt... vs. langsam, ruhig... (Erregung). Die Daten könnten dann mit Verfahren wie Distanz-Cluster-Analyse, Profilanalyse oder Faktorenanalyse ausgewertet werden.
- 7 Die Approximation der Prüfgröße χ^2_R an die χ^2 -Verteilung ist bereits bei verhältnismäßig kleinen Werten von n und k ($k = 3$, $n \geq 9$ bzw. $k = 4$, $n \geq 4$) relativ gut (vgl. FRIEDMAN 1940). Tabellen mit den exakten Wahrscheinlichkeiten gibt z.B. OWEN (1962:407ff).

- 8 Zur (quantitativen) Stilpsychologie (Stildiagnostik) vgl. z.B. MORIER (1959), SCHRÖDER (1960), SOMERS (1967) und GRAY (1969:20ff). SCHRÖDER untersucht insbesondere die Frage, "ob und inwieweit sich Intelligenz und Charakter auf quantitative Merkmale eines Textes auswirken" (1960:10). Zur Haltung GOETHEs gegenüber dem Verhältnis von Gestalt und Form vgl. die folgende Bemerkung aus "Dichtung und Wahrheit" (III,11): "Ich ehre den Rhythmus wie den Reim, wodurch Poesie erst zur Poesie wird, aber das eigentlich tief und gründlich Wirksame, das wahrhaft ausbildende und fördernde ist dasjenige, was vom Dichter übrig bleibt, wenn er in Prose übersetzt wird. Dann bleibt der reine, vollkommene Gehalt, den uns ein blendendes Äußere oft, wenn er fehlt, vorzuspiegeln weiß und wenn er gegenwärtig ist, verdeckt." (dtv-Gesamtausgabe 1962, Bd. 24, S. 47)
- 9 Da $E(p_i) = \sum p_i^2$ ist, kann R auch als Lokalisationsmaß interpretiert werden (vgl. ALTMANN/LEHFELDT 1979). Eine geometrische Interpretation von R findet sich bei HERDAN (1962:36). ONICESCU (1964) interpretiert $\sum p_i^2$ als Informationsenergie eines Systems. ONICESCU (1964) zeigt auch, daß die Informationsenergie abnimmt, wenn die Informationsentropie zunimmt. Zur Interpretation von R als Informationsenergie poetischer Texte vgl. MARCUS (1967).
- 10 Vgl. zum Folgenden vor allem MOOD (1940), DAVID/BARTON (1962:85ff, 119ff), WILKS (1962:144ff), BRUNK (1965:354ff), BROWNEE (1965:224ff), BRADLEY (1968:253ff), FISZ (1970:483ff). In der Metrik ist die Iterationstheorie von WORONCZAK (1961) verwendet worden.
- 11 Es gibt eine Reihe von Möglichkeiten zur Berechnung von $E(r)$ und $V(r)$. Neben der Methode der faktoriellen Momente erlaubt auch die Methode der Indikatorfunktionen eine relativ einfache Ableitung (vgl. z.B. MORAN 1968:38f, DAVID/BARTON 1962:85ff).
- 12 Vgl. WALLIS (1952). Näheres zur asymptotischen Verteilung der Iterationszahlen findet sich bei MOOD (1940) oder DAVID/BARTON (1962).
- 13 Vgl. auch LÜDTKE (1965), der bei der Ermittlung der Redundanz des lateinischen Hexameters, des altfranzösischen Epenverses, des französischen Alexandriners und des altgermanischen Stabreimverses zu dem Ergebnis kommt, daß diese in der Darstellung der konventionellen Metrik völlig heterogenen Systeme trotz großer Zeitdistanz und Unterschiede in der Sprachstruktur den gleichen Redundanzbetrag pro Silbe (0,6 bit) aufweisen. Ungefähr den gleichen Betrag hat LÜDTKE (1968) für das geographisch und sprachlich von den 1965 untersuchten Verssystemen stark abweichende "Epos von Herakleios" in Suahili ermittelt. Die Ergebnisse scheinen darauf hinzuweisen, daß die Redundanz ein "universaler ästhetischer Koeffizient" ist (LÜDTKE 1965:241).
- 14 Wie NEWMAN (1951) gezeigt hat, läßt sich auch die Autokorrelation von Vokalen und Konsonanten mit Hilfe periodischer Funktionen beschreiben.

- 15 Ähnlich wie der Index der Stärke der metrischen Bindung läßt sich auch ein Index der Reichweite der metrischen Bindung aufstellen. Aus beiden Indizes kann dann ein gemeinsamer Metrik-Index gebildet werden (vgl. FÜCKS 1968:67ff). Zur Charakterisierung der metrischen Bindung mit Hilfe der Entropie vgl. z.B. KONDRATOV (1969) und TURNER (1974).
- 16 Diese Aussage ist rein deskriptiv zu verstehen. Ob der Unterschied der Variationskoeffizienten statistisch signifikant ist, läßt sich natürlich auch testen. Ein Testverfahren gibt z.B. HALD (1967:301ff).
- 17 Vgl. auch das Ungarische, für das FÖNAGY (1961a:88f) - FÖNAGY stützt sich dabei auf TARNOCZY (1961) - folgende Werte angibt: wissenschaftliche Prosa: $\bar{x} = 2,5519$, Poesie: $\bar{x} = 1,9435$. Allerdings dürfte die wissenschaftliche Prosa einen extremen Wert darstellen. Weitere Belege finden sich bei GROTHJAHN (1977).
- 18 Vgl. auch MÜLLER (1964:41ff), der gezeigt hat, daß die mittlere Wortlänge in den Tragödien CORNEILLES signifikant größer ist als in den Komödien. Der Grund dürfte vor allem darin liegen, daß die Sprache der Komödien der Alltagssprache näher steht und somit frequentere Wörter aufweist.
- 19 Zur phonetischen Analyse von "Erkönig"-Rezitationen vgl. APPEL (1962) und WITTSACK (1962). APPEL hat die 1. Strophe des "Erkönig" u.a. nach Dauer, Tonhöhe und Intensität analysiert. WITTSACK hat die Intonationsverläufe ausgesuchter Verse als Ausdruck des "Naturmagischen" interpretiert.
- 20 Die Vernachlässigung des Stichprobenfehlers führt zu einer nicht unbeträchtlichen Rechnerleichterung und wirkt sich bei der gegebenen Stichprobengröße meist erst auf die 2. Stelle hinter dem Komma aus.
- 21 Vgl. auch MAAS (1971), der die Phonemhäufigkeiten in 12 Gedichtzyklen von Georg TRAKL untersucht hat. Wie die Daten von MAAS zeigen, kann auch bei TRAKL nicht pauschal von einer Verringerung der Frequenz der seltenen Phoneme und Erhöhung der Frequenz der häufigen Phoneme gesprochen werden.
- 22 Wie die Daten von MAAS (1971) und RYDER (1967) zeigen, bestehen in bezug auf das Verhältnis von Konsonant zu Vokal nicht unerhebliche Unterschiede zwischen verschiedenen poetischen Texten. RYDER hat die Konsonanten- bzw. Vokalfrequenz auch zu inhaltlichen Merkmalen von Gedichten in Beziehung gesetzt. Er kommt dabei zu dem Ergebnis, daß in "ruhigen" Gedichten wie GOETHE'S "Wanderers Nachtlied" im Gegensatz zu "bewegten" Gedichten wie GOETHE'S "An Schwager Kronos" die Konsonanten weniger frequent sind.
- 23 Ich habe bisher die Frage, inwieweit quantitative Textcharakteristiken als ästhetische Maßzahlen interpretiert werden können, ausgeklammert. Vgl. hierzu die Arbeiten von MASER (z.B. 1971) und vor allem von BENNE (z.B. 1969). Zur Kritik an der numerischen Ästhetik im Sinne von BENNE vgl. KRAUSE (1971).

13. Die generative Metrik

- 1 Etwa gleichzeitig mit HALLE/KEYSER (1966) ist HALLES wenig bekannte Veröffentlichung zur präislamischen Dichtung erschienen (siehe HALLE 1966), die ebenfalls zur generativen Metrik gerechnet werden kann. Als ein früher Vorläufer der generativen Metrik kann die "Axiomatik eines Verssystems, am mordwinischen Volkslied dargelegt" (1941) von JAKOBSON/LOTZ angesehen werden. Dieser Artikel beginnt mit folgender Feststellung: "Die Analyse eines metrischen Systems erfordert eine genaue Bestimmung aller Konstituenten und ihrer wechselseitigen Beziehungen, die jedem beliebigen Metrum dieses Systems zugrundeliegen; die Analyse muß vollständig und unmißverständlich klarstellen, welche Metren im gegebenen System vorkommen können und tatsächlich vorkommen, und welche nicht auftreten können. So muß der gesamte Bestand der wirklich vorhandenen metrischen Formen vollständig aus den aufgestellten Regeln ableitbar sein." BIERWISCH (1965:55) stellt hierzu fest: "Eine solche Versaxiomatik bildet eine genaue Parallele zur Grammatik G, sie ist ein Mechanismus, der alle möglichen Versmaße erzeugt und damit die Fähigkeit erklärt, Verse als solche zu verstehen, beabsichtigte oder irrtümliche Abweichungen zu erkennen usw." Sieht man einmal davon ab, daß BIERWISCH die aufgezeigte Parallele überinterpretiert und vergleicht die Feststellung BIERWISCHS mit Aussagen z.B. von HALLE/KEYSER (1966; 1971) zu den Zielen der generativen Metrik, so wird deutlich, daß BIERWISCH hier bereits das Programm der generativen Metrik formuliert hat. Es muß allerdings bezweifelt werden, ob dieses Programm überhaupt realisierbar ist. So hat auch LOTZ selbst kürzlich den sehr weitreichenden Anspruch der "Versaxiomatik" von 1941 revidiert: "The realization of a full corpus of metrical schemes and the exclusion of non-occurring and non-permitted metrical schemes is not possible due to the nature of verse itself, which, as the most idiosyncratic use of language, is subject to extreme individualistic manipulation. Therefore, instead of regarding the above account of Mordvinian verse as a closed explanatory system, it is better viewed as a descriptive statement which coherently defines the structural skeleton of a metric system..." (LOTZ 1974:980; vgl. auch LOTZ 1972).
- 2 Die generative Metrik ist jedoch nicht das erste Beispiel für die Anwendung generativer Methoden außerhalb der Grammatiktheorie. Bereits 1958, d.h. ein Jahr nach dem Erscheinen von CHOMSKY'S "Syntactic Structures", finden sich auf der "Conference on Style" (vgl. SEBEOK 1964) z.B. in den Beiträgen von VOEGELIN, SAPORTA und STANKIEWICZ Ansätze, generative Methoden innerhalb der Poetik zu verwenden. Der erste umfassendere Beitrag zur generativen Poetik wurde 1962 von LEVIN veröffentlicht. Als grundlegend für diese erste Phase der generativen Poetik ist OHMANN (1964) anzusehen. Das Hauptproblem dieser Arbeiten ist die Suche nach einem Regelmechanismus zur Beschreibung (Erklärung) syntaktischer und semantischer Abweichungen in der Poesie. Kontrovers ist vor allem der theoretische Status des Regelmechanismus. Zu den ver-

schiedenen Positionen (z.B. Subkomponente der Grammatik, unabhängige 'poetische Grammatik') vgl. z.B. THORNE (1965, 1969), BEZZEL (1969), BERGER (1972). Das Verhältnis der generativen Poetik zur generativen Metrik und der Stellenwert metrischer Regeln innerhalb der generativen (poetischen) Grammatik ist bisher erst vereinzelt diskutiert worden (vgl. z.B. VAN DIJK 1972, KIPARSKY 1972, IHWE 1975, R. FOWLER 1976). Diese Tatsache ist symptomatisch für das mangelnde Theoriebewußtsein der Vertreter der generativen Metrik, das u.a. auch von IHWE (1975) kritisiert wird.

- 3 Die folgenden Bemerkungen zur generativen Metrik beschränken sich auf einige wichtige Punkte. Einen Überblick über den Stand der Diskussion bis zum Erscheinen des von J.C. BEAVER und J.F. IHWE herausgegebenen Bandes "Generative Metrics" der Zeitschrift "Poetics" (1974) vermitteln die Artikel von KÜPER (1973) und BEAVER (1974).
- 4 Vgl. z.B. BEAVER (1968a,b, 1971a,b,c, 1974), FREEMAN (1968, 1969), HALLE (1970), HALLE/KEYSER (1971), HASCALL (1968, 1971), MAGNUSON/Ryder (1970, 1971). Bis Ende 1976 sind m.W. etwa 50 Artikel zur generativen Metrik erschienen. Eine umfassende, einführende Darstellung der generativen Metrik ist m.W. noch nicht veröffentlicht worden. Von deutschen Metrikern ist die generative Metrik bisher entweder überhaupt nicht oder äußerst kritisch rezipiert worden (zur Kritik vgl. etwa STANDOP 1972; KLEIN 1974; IHWE 1975). STANDOP meint sogar in bezug auf die Theorie von HALLE/KEYSER (1966, 1971), daß diese "auf dem Hintergrund der bisherigen metrischen Forschung nur als fataler Rückschritt gelten kann" (STANDOP 1972:1).
- 5 BEAVER (1974:7) unterscheidet drei Richtungen in der generativen Metrik: die "ursprüngliche Theorie" von HALLE/KEYSER (1966), die "erweiterte ursprüngliche Theorie" (z.B. BEAVER 1971a,b,c) und die "revidierte Theorie" von HALLE/KEYSER (1971).
- 6 KÜPER (1973:13) nennt u.a. folgende Parallelen zwischen generativer Transformationsgrammatik und generativer Metrik:
 - (i) Sowohl die generative Grammatik als auch die generative Metrik erheben den Anspruch, über eine formalisierbare, kohärente und explizite Theorie zu verfügen.
 - (ii) Beide Theorien erstreben neben deskriptiver auch explanative Adäquatheit.
 - (iii) Die generative Grammatik geht aus von der Sprachkompetenz des Hörer-Sprechers; die generative Metrik geht aus von der poetischen Kompetenz des Lesers.
 - (iv) Die generative Grammatik unterscheidet zwischen grammatischen und ungrammatischen Sätzen; die generative Metrik unterscheidet zwischen metrischen und unmetrischen Zeilen.
 - (v) Beide verwenden dazu Regeln, die alle grammatischen Sätze einer Sprache L bzw. alle metrischen Zeilen eines Metrums M generieren."

- 7 Vgl. auch HALLE (1970), der seine Theorie - allerdings in sehr knapper Form - an einigen nicht-englischen Metren exemplifiziert. Vgl. ferner die Ausführungen von ROUBAUD (1971) zum *verso de arte mayor* und von MAGNUSON (1974) zum slawischen Epenvers.
- 8 Was die Ziele der (generativen) Metrik betrifft, so vertritt z.B. STANDOP (1972:16) die Meinung, daß die Metrik nicht an der Untersuchung der metrischen Kompetenz, sondern primär an der Korpus-Analyse interessiert sei. IHWE (1975:379) ist der Ansicht, daß die Untersuchung der Beziehung zwischen metrischen Systemen und Sprachsystem die Hauptaufgabe einer generativen Metrik sei. In Kap. 3 und 6 habe ich dargelegt, daß die beiden genannten Ziele auch für die in dieser Arbeit vertretene Metrikkonzeption eine wichtige Rolle spielen.
- 9 Zu dem m.E. unbefriedigenden Versuch, stilistische Variation mit Hilfe des Parameters der Komplexität zu beschreiben vgl. BERNHART (1974).
- 10 Vgl. die Forderung von STANDOP (1972:17): "Regeln für eine genau umschriebene Versmenge müßten so angelegt sein, daß durch sie nicht einfach sämtliche vorkommenden Füllungstypen erzeugt würden - die häufigen ebenso wie die seltenen -, sie müßten vielmehr mit einem statistischen Parameter versehen sein, der die "relative Metrikalität" solcher Typen innerhalb des Korpus und damit so etwas wie ihren Akzeptabilitätsgrad festlegte."
- 11 Der *spondiacus* ist bei den meisten lateinischen Dichtern äußerst selten. So finden sich nach CRUSIUS/RUBENBAUER (1963:53) z.B. bei VERGIL lediglich etwa 0,24 % *spondiaci*. Dies gilt jedoch nicht in gleichem Maße für den griechischen Hexameter. Der Einfluß des griechischen Vorbildes zeigt sich z.B. bei CATULL im Epyllion, wo der Anteil der *spondiaci* nach CRUSIUS/RUBENBAUER (1963:53) etwa 6 % beträgt.
- 12 Vgl. BOLDORINI (1948:77), der den Einfluß von Lexikon und Morphologie als Erklärungsgrund anführt.
- 13 Auch HERDAN (1954) geht bei seiner informationstheoretischen Analyse von einem MARKOV-Modell aus. MARKOV-Ketten sind in der Linguistik z.B. bei der Untersuchung von Phonem- und Graphemfolgen (vgl. bereits MARKOV 1913) oder auch zur Bestimmung von Übergangswahrscheinlichkeiten zwischen grammatischen Kategorien (vgl. z.B. BRAINERD 1976) verwendet worden. Zur Kritik an der Verwendung des MARKOV-Modells im Bereich der Syntax vgl. CHOMSKY (1957:18ff).
- 14 Die folgenden Ausführungen zur MARKOV-Kette beruhen z.T. auf gemeinsamen Überlegungen mit D. JAESCHKE vom Institut für Statistik der Universität Dortmund.
- 15 Bei einer probabilistischen Bewertung der Transformationsregeln von DEVINE/STEPHENS (1975) müßte allerdings zusätzlich die identische Transformation $\sim \sim \sim \rightarrow \sim \sim \sim$ eingeführt werden.

- 16 So spielt es z.B. für eine Computeranalyse metrischer Texte keine Rolle, ob das für die Analyse verwendete Regelsystem psychische Vorgänge widerspiegelt oder nicht. Deshalb sind auch die von DEVINE/STEPHENS (1975) vorgebrachten psychologischen Argumente nur dann valide, wenn man für die generative Metrik eine entsprechende Zielsetzung postuliert. Zielsetzungen bleiben jedoch stets 'Setzungen' und sind erfahrungswissenschaftlich nicht begründbar. Man sollte sich aus diesem Grunde auch vor allzu apodiktischen Feststellungen etwa folgender Art hüten: "The questions that generative grammar seeks to answer are fundamentally psychological questions..." (KIPARSKY 1972:171).
- 17 Diese Einschränkung hätte man auch von vornherein z.B. durch folgende Spezifizierung der Regel R_{ij} erreichen können:
 $A_1 = d$, für $j \neq 6$
 $A_3 = t$, für $j = 6$
- 18 Bei der probabilistischen Bewertung von Grammatiken wird meist davon ausgegangen, daß zwischen den Regeln keine Abhängigkeiten bestehen (vgl. z.B. SUPPES 1972). Bedingte Wahrscheinlichkeiten werden bisher erst vereinzelt verwendet (vgl. z.B. SALOMAA 1969; KLEIN 1974a:109ff).

A N H A N G

ERLKÖNIG

- I 1 Wer reitet so spät durch Nacht und Wind?
 2 Es ist der Vater mit seinem Kind;
 3 Er hat den Knaben wohl in dem Arm,
 4 Er faßt ihn sicher, er hält ihn warm.—
- II 5 Mein Sohn, was birgst du so bang dein Gesicht?—
 6 Siehst, Vater, du den Erlkönig nicht?
 7 Den Erlenkönig mit Kron' und Schweif?—
 8 Mein Sohn, es ist ein Nebelstreif.—
- III 9 "Du liebes Kind, komm, geh mit mir!
 10 Gar schöne Spiele spiel' ich mit dir,
 11 Manch' bunte Blumen sind an dem Strand;
 12 Meine Mutter hat manch' gülden Gewand."
- IV 13 Mein Vater, mein Vater, und hörest du nicht,
 14 Was Erlenkönig mir leise verspricht?—
 15 Sei ruhig, bleibe ruhig, mein Kind!
 16 In dünnen Blättern säuselt der Wind.—
- V 17 "Willst, feiner Knabe, du mit mir gehn?
 18 Meine Töchter sollen dich warten schön;
 19 Meine Töchter führen den nächtlichen Reihn
 20 Und wiegen und tanzen und singen dich ein."
- VI 21 Mein Vater, mein Vater, und siehst du nicht dort
 22 Erlkönigs Töchter am düstern Ort?—
 23 Mein Sohn, mein Sohn, ich seh' es genau;
 24 Es scheinen die alten Weiden so grau.—

- VII 25 "Ich liebe dich, mich reizt deine schöne Gestalt,
 26 Und bist du nicht willig, so brauch' ich Gewalt."—
 27 Mein Vater, mein Vater, jetzt faßt er mich an!
 28 Erlkönig hat mir ein Leids getan!—
- VIII 29 Dem Vater grauset's, er reitet geschwind,
 30 Er hält in Armen das ächzende Kind,
 31 Erreicht den Hof mit Mühe und Not,
 32 In seinen Armen das Kind war tot.

(zit. nach der Ausgabe von TRUNZ 1962)

LITERATURVERZEICHNIS

- ABERCROMBIE, David (1964): A phonetician's view of verse structure. In: *Linguistics* 6, 5-13
- ABERNATHY, Robert (1967): Rhymes, non-rhymes, and antirhymes. In: *To Honor Roman Jakobson I, Den Haag*, 1-14
- ABRAHAM, Werner (1971): Stil, Pragmatik und Abweichungsgrammatik. In: A. v. STECHOW (Hg.): *Beiträge zur generativen Grammatik*. Braunschweig, 1-13
- ABRAHAM, W./BRAUNMÜLLER, K. (1971): Stil, Metapher und Pragmatik. In: *Lingua* 28, 1-47
- ADORNO, Theodor W. u.a. (1974): *Der Positivismusstreit in der deutschen Soziologie*. Darmstadt/Neuwied 1969
- ALBERT, Hans (1968): *Traktat über kritische Vernunft*. Tübingen
- ALBERT, Hans (1957): Theorie und Prognose in den Sozialwissenschaften. In: *Schweizerische Zeitschrift für Volkswirtschaft und Statistik* 93, 60-76 (abgedruckt in: Ernst TOPITSCH (Hg.): *Logik der Sozialwissenschaften*. Köln 1972, 126-143)
- ALBERT, Hans/KEUTH, Herbert (1973): *Kritik der kritischen Psychologie*. Hamburg
- ALBRECHT, JÖRN (1973): *Linguistik und Übersetzung*. Tübingen (=Romanistische Arbeitshefte 4)
- ALBRIGHT, W.F. (1950): The Psalm of Habakkuk. In: *Studies presented to T. H. ROBINSON*. Edinburgh, 1-18
- ALDEN, Raymond M. (1914): The Mental Side of Metrical Form. In: *Modern Language Review* 9, 297-308
- ALTMANN, Gabriel (1973): Mathematische Linguistik. In: KOCH (1973), 208-232
- ALTMANN, Gabriel (1972): Status und Ziele der quantitativen Sprachwissenschaft. In: S. JÄGER (Hg.): *Linguistik und Statistik*. Braunschweig, 1-9
- ALTMANN, Gabriel (1971): Rezension zu L. DOLEŽEL/R.W. BAILEY (Hg.): *Statistics and Style*. New York 1969. In: *Muttersprache* 4(1971), 276-282
- ALTMANN, Gabriel (1971a): *Introduction to Quantitative Phonology*. Bochum: Habilschrift
- ALTMANN, Gabriel (1969): Differences between Phonemes. In: *Phonetica* 19, 118-132

- ALTMANN, Gabriel (1966a): Binomial Index of Euphony for Indonesian Poetry. In: Asian and African Studies 2, 62-67
- ALTMANN, Gabriel (1966b): The measurement of euphony. In: Theorie Verses I, 263-264
- ALTMANN, Gabriel (1963): Phonic Structure of Malay Pantun. In: Archiv Orientální 31, 274-286
- ALTMANN, Gabriel/LEHFELDT, Werner (1979): Einführung in die quantitative Phonologie. Konstanz (im Druck)
- ALTMANN, Gabriel/LEHFELDT, Werner (1973): Allgemeine Sprachtypologie. München
- ALTMANN, Gabriel/ŠTUKOVSKÝ, Robert (1965): The Climax in Malay Pantun. In: Asian and African Studies 1, 13-20
- ANDRESEN, Helga (1974): Der Erklärungsgehalt linguistischer Theorien. München
- APPEL, Wilhelm (1962): Gestaltstudien an deutschen Versen. In: Zeitschrift für Phonetik, Sprachwissenschaft und Kommunikationsforschung 15, 21-38
- ARENS, Hans (1965): Verborgene Ordnung. Die Beziehungen zwischen Satzlänge und Wortlänge in deutscher Erzählprosa vom Barock bis heute. Düsseldorf
- ARMITAGE, P. (1955): Tests for linear trends in proportions and frequencies. In: Biometrics 11, 375-386
- ARNOLD, Heinz L./SINEMUS, Volker (1976): Grundzüge der Literatur- und Sprachwissenschaft. Bd. 1: Literaturwissenschaft. Hrsg. von H. L. ARNOLD/V. SINEMUS, München 1973
- ASMUTH, Bernhard (1976): Klang - Metrum - Rhythmus. In: ARNOLD/SINEMUS (1976), 208-227
- AUSTERLITZ, R. (1958): Ob-Ugric Metrics. The metrical structure of Ostyak and Vogul folk poetry. Helsinki (=Folklore fellows communications 174)
- AVALLE, D'Arco Silvio (1963): Preistoria dell'Endecasillabo. Milano/Napoli
- BAEHR, Rudolf (1962): Spanische Verslehre auf historischer Grundlage. Tübingen
- BAEHR, Rudolf (1970): Einführung in die französische Verslehre. München
- BAILEY, James (1975): Linguistic Givens and their Metrical Realization in a Poem by Yeats. In: Language and Style 8, 21-35
- BAILEY, James (1973): Towards a Statistical Analysis of English Verse. The iambic tetrameter of ten poets. The Hague/Paris

- BALDWIN, J. R. (1969): Syllable Division in Russian. In: Zeitschrift für Phonetik, Sprachwissenschaft und Kommunikationsforschung 22, 211-217
- BARTON, D.E./DAVID, F. N. (1957): Multiple Runs. In: Biometrika 44, 168-178 (Corrigenda S. 534)
- BARTSCH, Renate/VENNEMANN, Theo (1973): Sprachtheorie. In: H.P. ALTHAUS/H. HENNE/H. E. WIEGAND (Hg.): Lexikon der Germanistischen Linguistik. Studienausgabe Bd. I. Tübingen, 34-55
- BASHARIN, G. P. (1959): On a Statistical Estimate for the Entropy of a Sequence of Independent Random Variables. In: Theory of Probability and its Application 4 (1959), 333-336
- BEAVER, Joseph C. (1974): Generative Metrics: The Present Outlook. In: Poetics 12, 7-28
- BEAVER, Joseph C. (1971a): Current Metrical Issues. In: College English 33, 177-197
- BEAVER, Joseph C. (1971b): The rules of stress in English verse. In: Language 47, 586-614
- BEAVER, Joseph C. (1971c): Review of Halle and Keyser 1971. In: Language Sciences 18, 33-38
- BEAVER, Joseph C. (1968a): Progress and Problems in Generative Metrics. In: B. J. DARDEN u.a. (Hg.): Papers from the Fourth Regional Meeting, Chicago Linguistic Society, April 19-20, 1968, Chicago, 146-155
- BEAVER, Joseph C. (1968b): A grammar of prosody. In: College English 29, 310-321 (abgedruckt in: FREEMAN (1970), 427-448)
- BECHERT, Johannes (1971): Die Natürlichkeit linguistischer Theorien. In: Linguistische Berichte 13, 49-54
- BECK, Miroslav (1966): Einige Bemerkungen zur vergleichenden Metrik. In: Theorie Verses I (1966), 81-85
- BELLYN, D. E. (1975): Studies in the New Experimental Aesthetics. New York/London/Sidney
- BELYJ, Andrej (1910): Simvolizm. Moskva
- BENARY, Peter (1967): Rhythmik und Metrik. Köln
- BENEŠ, Eduard/VACHEK, Joseph (1971): Stilistik und Soziolinguistik. Beiträge der Prager Schule zur strukturellen Sprachbetrachtung und Spracherziehung. Hrsg. von E. BENEŠ/J. VACHEK. Berlin
- BENNINGHAUS, Hans (1974): Deskriptive Statistik. Stuttgart

- BENSE, Max (1969): Einführung in die informationstheoretische Ästhetik. Reinbek
- BENVISTE, Emile (1951): La notion de 'rythme' dans son expression linguistique. In: Journal de Psychologie normale et pathologique 44, 401-410
- BENUSSI, Vittorio (1913): Psychologie der Zeitauffassung. Heidelberg
- BERGER, Albert (1972): Poesie zwischen Linguistik und Literaturwissenschaft. In: Linguistische Berichte 17, 1-11
- BERNHART, Walter A. (1974): Complexity and Metricality. In: Poetics 12, 113-141
- BERSCHIN, Helmut (1970): Automatische metrische Analyse lateinischer Verse. In: Linguistik und Didaktik 1, 72-80
- BEUCHOT, M. (1830): Arouet de Voltaire: Oeuvres, Bd. II. Hrsg. von M. BEUCHOT. Paris
- BEZZEL, Chris (1969): Some Problems of a Grammar of Modern German Poetry. In: Foundations of Language 5, 470-487
- BIERWISCH, Manfred (1966): Strukturalismus. In: Kursbuch 5, 77-152
- BIERWISCH, Manfred (1965): Poetik und Linguistik. In: Mathematik und Dichtung (1965), 49-66
- BIERWISCH, Manfred/HEIDOLPH, Karl Erich (1970): Progress in Linguistics. Hrsg. von M. BIERWISCH/K. E. HEIDOLPH. The Hague 1968
- BLOCH, Bernhard (1953): Linguistic Structure and Linguistic Analysis. In: A. A. HILL (Hg.): Report of the Fourth Annual Round Table Meeting on Linguistics and Language Teaching. Washington, D.C., 40-44
- BLÖSCHL, Lilian (1966): Kullbacks 2 $\hat{I}$ -Test als ökonomische Alternative zur Chi²-Probe. In: Psychologische Beiträge 9, 379-406
- BOLDRINI, Marcello (1948): Esametri. In: Contributi del Laboratorio di Statistica, Serie Sesta, dell' Università Cattolica del Sacro Cuore 21, 76-79
- BOLTON, Thaddeus L. (1894): Rhythm. In: The American Journal of Psychology 6, 145-238
- BOUDON, Raymond (1972): Mathematische Modelle und Methoden. Frankfurt
- BOWLEY, C.C. (1974): Metrics and the generative approach. In: Linguistics 121, 5-19
- BOYD, James (1961): Notes to Goethe's Poems. Vol. I. Oxford
- BRADLEY, James V. (1968): Distribution-Free Statistical Tests. Englewood Cliffs, N.J.

- BRAINERD, Barron (1976): On the Markov Nature of Text. In: Linguistics 176, 5-30
- BREUER, Dieter (1974): Einführung in die pragmatische Texttheorie. München
- BROWNLEE, K.A. (1965): Statistical Theory and Methodology in Science and Engineering. New York/London
- BRÜCKE, E. (1871): Die physiologischen Grundlagen der neuhochdeutschen Verskunst. Wien
- BRUNER, J.S./GOODNOW, J./AUSTIN, G. A. (1956): A Study of Thinking. New York
- BRUNK, H. D. (1965): An Introduction to Mathematical Statistics. Waltham, Mass.
- BÜHLER, Karl (1918): Kritische Musterung der neueren Theorien des Satzes. In: Indogermanisches Jahrbuch 6, 1-20
- BÜHLER, Karl (1965): Sprachtheorie. Die Darstellungsfunktion der Sprache. Stuttgart (1934)
- BUNGE, Mario (1967): Scientific Research. Vol. I/II. Berlin/Heidelberg/New York
- BURNET, Joannes (1962): Platonis opera. Tomus V. ed. J. BRUNET. Oxford
- CAMARA, J. Mattoso Jr. (1953): Para o estudo da fonêmica portuguesa. Rio de Janeiro
- CARLOS, Federico/ROBLES, Sainz de (1961): Lope Felix de Vega Cappio: Obras escogidas. Bd. II. Hrsg. von F. CARLOS/S. de Robles. Madrid
- CEDERGREN, Henrietta/SANKOFF, David (1974): Variable Rules: Performance as a Statistical Reflection of Competence. In: Language 50, 333-355
- ČERVENKA, Miroslav (1973): Die Grundkategorien des Prager literaturwissenschaftlichen Strukturalismus. In: V. ŽMEGAC/Z. SKREB: Zur Kritik literaturwissenschaftlicher Methodologie. Frankfurt, 137-168
- ČERVENKA, Miroslav (1971): Die Umgestaltungen des tschechischen Alexandriners. In: Zeitschrift für Literaturwissenschaft und Linguistik 3, 107-124
- ČERVENKA, Miroslav/SGALLOVÁ, Květa (1967): On a probabilistic model of the Czech Verse. In: Prague Studies in Mathematical Linguistics 2, 105-120

- CHASTAING, Maxime (1964): Nouvelles recherches sur le symbolisme des voyelles. In: *Journal de psychologie et de pathologie* 1, 75-88
- CHATMAN, Seymour (1965): *A Theory of Meter*. Den Haag
- CHATMAN, Seymour (1964): Comparing Metrical Styles. In: SEBEDK (1964), 149-172
- CHERRY, Collin (1967): *Kommunikationsforschung - eine neue Wissenschaft*. Hamburg
- CHOMSKY, Noam (1965): *Aspects of the Theory of Syntax*. Cambridge, Mass.
- CHOMSKY, Noam (1957): *Syntactic Structures*. The Hague
- CHOMSKY, Noam/HALLE, Morris (1968): *The Sound Pattern of English*. New York
- CLARKE, Dorothy C. (1937): *Una bibliografía de versificación española*. Berkeley
- COCHRAN, William G. (1963): *Sampling Techniques*. New York
- COCHRAN, William G. (1954): Some Methods for Strengthening the Common χ^2 Tests. In: *Biometrics* 10, 417-451
- COHEN, Jean (1966): *Structure du langage poétique*. Paris
- COLLINDER, Björn (1951): Three Degrees of Quantity. In: *Studia Linguistica* 5, 28-43
- CONSRUCH, Maximilian (1906): *Hephaestionis Enchiridion*. Leipzig
- CORCORAN, D. W. J. (1971): *Pattern Recognition*. Harmondsworth
- COSERIU, Eugenio (1952): *Sistema, norma y habla*. Montevideo (deutsche Übersetzung in: E. COSERIU: *Sprachtheorie und Allgemeine Sprachwissenschaft*. München 1975, 11-101)
- COSTNER, Herbert L. (1965): Criteria for Measures of Association. In: *American Sociological Review* 30, 341-353
- CROSS, F. M. Jr. (1950): Notes on a Canaanite Psalm in the Old Testament. In: *Bulletin of the American Schools of Oriental Research* 117, 19-21
- CRUSIUS, Friedrich/RUBENBAUER, Hans (1963): *Römische Metrik*. München
- CZERNY, Zygmunt (1961): Le vers libre et son art structural. In: *Poetics, Poetyka, Poëtika I* (1961), 249-279
- DAVID, F. N./BARTON, D. E. (1962): *Combinatorial Chance*. London
- DELATTRE, Pierre (1965): Comparing the phonetic features of English, French, German and Spanish. Heidelberg

- DELBOUILLE, Paul (1967): Recherches récentes sur la valeur suggestive des sonorités. In: PARENT (1967). 141-162
- DELBOUILLE, Paul (1961): *Poésie et Sonorités*. Bruxelles
- DEVINE, Andrew M./STEPHENS, Laurence D. (1975): The Abstractness of Metrical Patterns: Generative Metrics and Explicit Traditional Metrics. In: *Poetics* 16, 411-429
- DIJK, Teun A. van (1972): Some Aspects of Text Grammars. A Study in Theoretical Linguistics and Poetics. The Hague
- DIJK, Teun A. van (1972a): *Beiträge zur generativen Poetik*. München
- DOLEŽEL, Lubomír (1969): A Framework for the Statistical Analysis of Style. In: L. DOLEŽEL/R. W. BAILEY: *Statistics and Style*. New York, 10-25 (deutsche Übersetzung in: IHWE (1971/72) Bd. I, 253-273)
- DOLEŽEL, Lubomír (1967): The Prague School and the Statistical Theory of Poetic Language. In: *Prague Studies in Mathematical Linguistics* 2, 97-104
- DOLEŽEL, Lubomír (1965): Zur statistischen Theorie der Dichtersprache. In: *Mathematik und Dichtung* (1965), 275-293
- DOLEŽEL, Lubomír (1963): Prědběžný odhad entropie a redundance psané češtiny. In: *Slovo a slovestnost* 24, 165-175
- DOLEŽEL, Lubomír/BAILEY, Richard W. (1969): *Statistics and Style*. Hrsg. von L. DOLEŽEL/R. W. BAILEY. New York
- DROBISCH, W. M. (1866): Ein statistischer Versuch über die Formen des lateinischen Hexameters. In: *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Philologisch-historische Klasse, Band 18*. Leipzig, 75-139
- DURAND, Marguerite (1954): La syllabe. Ses définitions. Sa nature. In: *Orbis* 3, 527-533
- EIBL, Karl (1976): *Kritisch-rationale Literaturwissenschaft*. München
- EIMERMACHER, Karl (1969): Entwicklung, Charakter und Probleme des sowjetischen Strukturalismus in der Literaturwissenschaft. In: *Sprache im Technischen Zeitalter* 30, 126-156
- ELWERT, Wilhelm Theodor (1968): *Italienische Metrik*. München
- ELWERT, Wilhelm Theodor (1966): *Französische Metrik*. München 1966
- ENGLER, Rudolf (1962): Théorie et critique d'un principe saussurien: l'arbitraire du signe. In: *Cahiers Ferdinand de Saussure* 19, 5-66
- ENGLER, Rudolf (1964): Compléments à l'arbitraire. In: *Cahiers Ferdinand de Saussure* 21, 25-32

- ENKVIST, Nils Erik (1973): *Linguistic Stylistics*. The Hague/Paris
- ERLICH, Victor (1965): *Russian Formalism*. The Hague
- ERN, L. (1968): *Freivers und Metrik. Zur Problematik der englischen Verswissenschaft*. Diss. Freiburg
- ERTEL, Suitbert (1969): *Psychophonetik. Untersuchungen über Lautsymbolik und Motivation*. Göttingen
- ESSEN, Otto von (1966): *Allgemeine und angewandte Phonetik*. Berlin
- ESSEN, Otto von (1961): *Mathematische Analyse periodischer Vorgänge in gemeinfaßlicher Darstellung*. Marburg
- ESSEN, Otto von (1951): *Die Silbe - ein phonologischer Begriff*. In: *Zeitschrift für Phonetik* 5, 199-203
- ESSLER, Wilhelm K. (1971): *Wissenschaftstheorie II. Theorie und Erfahrung*. Freiburg/München
- ESSLER, Wilhelm K. (1970): *Wissenschaftstheorie I. Definition und Reduktion*. Freiburg/München
- FAULSTICH, Werner (1976): *Die Relevanz der Cloze-Procedure als Methode wissenschaftlicher Textuntersuchung. Ein Beitrag zur Literaturwissenschaft als Sozialwissenschaft*. In: *Zeitschrift für Literaturwissenschaft und Linguistik* 21, 80-95
- FAURE, Georges (1970): *Les éléments du rythme poétique en anglais moderne: Esquisse d'une nouvelle analyse et essai d'application au Prométhée Unbound de P. B. Shelley*. The Hague/Paris
- FAURE, Georges (1964): *Le rôle du rendement fonctionnel dans la perception des oppositions vocaliques distinctives du français*. In: *In Honour of Daniel JONES*. London, 320-328
- FEYERABEND, Paul K. (1975): *Against Method. Outline of an anarchistic theory of knowledge*. London
- FIGGE, Udo L. (1973): *Strukturelle Linguistik*. In: KOCH (1973), 1-36
- FINAERT, Guy/THONNARD, F.J. (1947): *Oeuvres de Saint Augustin, t. 7*. Hrsg. von G. FINAERT/F.J. THONNARD. Bruges
- FISCHER, Walther L. (1976): *Mathematische Texttheorie*. In: ARNOLD/SINEMUS (1976), 44-60
- FISCHER, Walther L. (1973): *Äquivalenz- und Toleranzstrukturen in der Linguistik. Zur Theorie der Synonyma*. München
- FISCHER, Walther L. (1972): *Texte und Zufallsfolgen*. In: H. SCHANZE (Hg.): *Literatur und Datenverarbeitung*. Tübingen, 192-209
- FISCHER, Walther L. (1970): *Beispiele topologischer Stilcharakteristiken*. In: *Grundlagenstudien aus Kybernetik und Geisteswissenschaften* 11, 1-11

- FISCHER, Walther L. (1970a): *Automorphismengruppen von Texten. Ein Beitrag zur Algebra der Texte*. In: E. WALTHER/L. HARIG (Hg.): *Muster möglicher Welten*. Wiesbaden, 41-45
- FISCHER, Walther L. (1970b): *Mathematik und Literaturtheorie*. In: *Sprache im technischen Zeitalter* 34, 106-120
- FISCHER, Walther L. (1969): *Texte als simpliziale Komplexe*. In: *Beiträge zur Linguistik und Informationsverarbeitung* 17, 27-48
- FISCHER-JØRGENSEN, E. (1952): *On the definition of phoneme categories*. In: *Acta Linguistica* 7, 8-39
- FISZ, Marek (1970): *Wahrscheinlichkeitsrechnung und mathematische Statistik*. Berlin (polnisches Original 1967)
- FLETCHER, Harvey (1965): *Speech and Hearing in Communication*. New York 1953
- FÓNAGY, Ivan (1971): *Le signe conventionnel motivé*. In: *La linguistique* 7, 55-80
- FÓNAGY, Ivan (1971a): *The Functions of Vocal Style*. In: S. CHATMAN (Hg.): *Literary Style: A Symposium*. London/New York, 159-176
- FÓNAGY, Ivan (1965): *Der Ausdruck als Inhalt*. In: *Mathematik und Dichtung* (1965), 243-274
- FÓNAGY, Ivan (1963): *Die Metaphern in der Phonetik*. Den Haag
- FÓNAGY, Ivan (1961): *Informationsgehalt von Wort und Laut in der Dichtung*. In: *Poetics, Poetyka, Poëtika I* (1961), 591-605
- FÓNAGY, Ivan (1961a): *Die Silbenzahl der ungarischen Wörter in der Rede*. In: *Zeitschrift für Phonetik* 14, 88-92
- FOUCHÉ, Pierre (1930): *La théorie dynamique de la syllabe*. In: *Bericht über die 1. Tagung der Internationalen Gesellschaft für experimentelle Phonetik in Bonn*. Bonn, 32-35
- FOWLER, Roger (1966): *Structural Metrics*. In: *Linguistics* 27, 49-64
- FOWLER, Roger (1966a): *'Prose Rhythm' and Metre*. In: R. FOWLER (Hg.): *Essays on Style and Language*. London, 82-99
- FOWLER, Rowena (1976): *Metrics and the Transformational-Generative Model*. In: *Lingua* 38, 21-36
- FRAISSE, Paul (1963): *The Psychology of Time*. New York
- FRAISSE, Paul (1956): *Les structures rythmiques*. Paris
- FREEMAN, Donald C. (1970): *Linguistics and literary style*. Hrsg. von D. C. FREEMAN. New York

- FREEMAN, Donald C. (1969): Metrical position constituency and generative metrics. In: *Language and Style* 2, 195-206
- FREEMAN, Donald C. (1968): On the Primes of Metrical Style. In: *Language and Style* 1, 63-101 (abgedruckt in: FREEMAN (1970), 448-491)
- FREY, Gerhard (1967): *Die Mathematisierung unserer Welt*. Stuttgart
- FRIEDMAN, Milton (1940): A Comparison of Alternative Tests of Significance for the Problem of m rankings. In: *Annals of Mathematical Statistics* 11, 86-92
- FRIEDMAN, Milton (1937): The use of ranks to avoid the assumption of normality implicit in the analysis of variance. In: *Journal of the American Statistical Association* 32, 675-701
- FRIEDRICH, Walter/HENNIG, Werner (1975): *Der sozialwissenschaftliche Forschungsprozeß*. Hrsg. von W. FRIEDRICH/W. HENNIG. Berlin
- FRIEDRICHS, Jürgen (1973): *Methoden empirischer Sozialforschung*. Reinbek
- FROMKIN, Victoria A. (1968): Speculations on performance models. In: *Journal of Linguistics* 4, 47-68
- FRY, Denis B. (1970): Prosodic Phenomena. In: MALMBERG (1970): 365-410
- FRY, Denis B. (1971): Speech Reception and Perception. In: J. LYONS (Hg.): *New Horizons in Linguistics*. Harmondsworth 1970, 29-52
- FUCKS, Wilhelm (1971): Possibilities of Exact Style Analysis. In: J. STRELKA (Hg.): *Patterns of Literary Style*. Pennsylvania State University, 51-75
- FUCKS, Wilhelm (1968): *Nach allen Regeln der Kunst*. Stuttgart
- FUCKS, Wilhelm (1963): Mathematische Analyse von Formalstrukturen von Werken der Musik. In: *Arbeitsgemeinschaft für Forschung des Landes Nordrhein-Westfalen* 124, 39-93
- FUCKS, Wilhelm (1955): *Mathematische Analyse von Sprachelementen, Sprachstil und Sprachen*. Köln/Opladen
- FUCKS, Wilhelm (1953): Mathematische Analyse des literarischen Stils. In: *Studium Generale* 6, 506-523
- FUCKS, Wilhelm/LAUTER, Josef (1965): Mathematische Analyse des literarischen Stils. In: *Mathematik und Dichtung* (1965), 107-122
- FUCKS, Wilhelm/LAUTER, Josef (1965a): *Exaktwissenschaftliche Musikanalyse*. Köln/Opladen 1965 (Forschungsberichte NRW Nr. 1519)
- FUDGE, Erik (1969): Syllables. In: *Journal of Linguistics* 5, 253-286
- GÁLDI, Ladislav (1967): Le mètre et ses variantes typiques. In: To honor R. Jakobson I. The Hague/Paris, 692-696

- GAMES, Paul A./WINKLER, Henry B./PROBERT, David A. (1972): Robust Tests for Homogeneity of Variance. In: *Educational and Psychological Measurement* 32(1972), 887-909
- GARVIN, Paul L. (1964): *A Prague School Reader on Esthetics, Literary Structure and Style*. Hrsg. von P. L. GARVIN. Washington
- GARY-PRIEUR, Marie-Noëlle (1971): La notion de connotation(s). In: *Littérature* 4, 96-107
- GEBHARDT, Friedrich (1966): Verteilung und Signifikanzschranken des 3. und 4. Stichprobenmoments bei normalverteilten Variablen. In: *Biometrische Zeitschrift* 8, 219-241
- GHISELLI, Edwin E. (1964): *Theory of Psychological Measurement*. New York
- GIESEN, Bernhard/SCHMID, Michael (1975): *Theorie, Handeln und Geschichte. Erklärungsprobleme in den Sozialwissenschaften*. Hrsg. von B. GIESEN/M. SCHMID. Hamburg
- GIGON, Olof (1950): *Aristoteles: Vom Himmel. Von der Seele. Von der Dichtkunst*. Übersetzt von O. GIGON. Zürich
- GOLDENBERG, Yves (1976): La métrique arabe classique et la typologie métrique. In: *Revue Roumaine de Linguistique - Cahiers de linguistique théorique et appliquée* 13, 85-98
- GONDA, J. (1959): *Stylistic Repetition in the Veda*. Amsterdam
- GOODMAN, Leo A./KRUSKAL, William H. (1954): Measures of Association for Cross Classifications. In: *Journal of the American Statistical Association* 49, 732-764
- GORDON, Cyrus (1965): *Ugaritic Textbook*. Rom (=Analecta Orientalia 38)
- GRAMMONT, Maurice (1923): *Le vers français*. Paris
- GRANDA GUTIERREZ, Germán de (1966): La estructura silábica y su influencia en la evolución fonética del dominio ibero-románico. Madrid
- GRAUBNER, Hans (1976): Stilistik. In: ARNOLD/SINEMUS (1976), 164-187
- GRAY, Benisson (1969): *Style. The Problem and its Solution*. The Hague/Paris
- GREGER, Karl (1972): Zufälliges Zusammenklumpen. In: *Aus der Schulpraxis - Für die Schulpraxis* 25, 211-214
- GROEBEN, Norbert (1976): Empirische Literaturwissenschaft als Metatheorie. In: *Zeitschrift für Literaturwissenschaft und Linguistik* 21, 125-145
- GROEBEN, Norbert (1972): *Literaturpsychologie. Literaturwissenschaft zwischen Hermeneutik und Empirie*. Stuttgart
- GROOT, A. Willem de (1946): *Algemene Versleer*. Den Haag

- GROOT, A. Willem de (1932): Der Rhythmus. In: *Neophilologus* 17, 81-100, 177-197, 241-265
- GROOT, A. Willem de (1918): A Handbook of Antique Prose-Rhythm. Groningen
- GROSSE, Ernst Ulrich (1971): *Altfranzösischer Elementarkurs*. München
- GROTJAHN, Rüdiger (1977): Die Verteilung der Anzahl der Silben pro Wort. (Manuskript)
- GUÉRON, Jacqueline (1974): The Meter of Nursery Rhymes. An Application of the Halle-Keyser Theory of Meter. In: *Poetics* 12, 73-111
- GUEUNIER, Nicole (1969): La pertinence de la notion d'écart en stylistique. In: *Langue française* 3, 34-45
- GUIRAUD, Pierre (1970): *La versification*. Paris
- GUIRAUD, Pierre (1953): *Langage et versification d'après l'oeuvre de Paul Valéry*. Paris
- GUTJAHR, Walter (1971): *Die Messung psychischer Eigenschaften*. Berlin
- HABERMAS, Jürgen (1973): Wahrheitstheorien. In: *Wirklichkeit und Reflexion*. Walter SCHULZ zum 60. Geburtstag. Hrg. von H. FAHRENBACH. Pfullingen, 211-265
- HÁLA, Bohuslav (1973): La sílaba, su naturaleza, su origen y sus transformaciones. Madrid 1966
- HALD, A. (1967): *Statistical Theory with Engineering Applications*. New York/London
- HALLE, Morris (1970): On Meter and Prosody. In: BIERWISCH/HEIDOLPH (1970), 64-80
- HALLE, Morris (1966): On the metrics of pre-Islamic poetry. In: MIT - Research Laboratory of Electronics, Quaterly Progress Report 83, 113-116
- HALLE, Morris/KEYSER, Samuel J. (1972): English III. The Iambic Pentameter. In: WIMSATT (1972), 217-237 (gekürzte Fassung von HALLE/KEYSER 1971, Kap. 3)
- HALLE, Morris/KEYSER, Samuel J. (1971): *English Stress: Its Form, its Growth, and its Role in Verse*. New York
- HALLE, Morris/KEYSER, Samuel J. (1966): Chaucer and the study of prosody. In: *College English* 28, 187-219 (abgedruckt in: FREEMAN (1970), 366-426)
- HAMMARSTRÖM, Göran (1966): *Linguistische Einheiten im Rahmen der modernen Sprachwissenschaft*. Berlin/Heidelberg/New York
- HANSON, Göte (1967): *Dimensions in speech sound perception. An experimental study of vowel perception*. Uppsala

- HARWEG, Roland (1968): *Pronomina und Textkonstitution*. München
- HARWEG, R./SUERBAUM, U./BECKER, H. (1967): Sprache und Musik. In: *Poetica* 1, 390-414
- HASCALL, Dudley L. (1974): Triple Meter in English Verse. In: *Poetics* 12, 49-71
- HASCALL, Dudley L. (1971): Trochaic Meter. In: *College English* 33, 217-226
- HASCALL, Dudley L. (1968): Some Contributions to the Halle-Keyser Theory of Prosody. In: *College English* 30, 357-365
- HAUGEN, Einar (1956): The Syllable in Linguistic Description. In: For Roman Jakobson. The Hague, 213-221
- HAVRÁNEK, Bohuslav (1964): The Functional Differentiation of the Standard Language. In: GARVIN (1964), 3-16
- HAYS, William L. (1973): *Statistics for the Social Sciences*. London/New York
- HEIKE, Georg (1972): *Phonologie*. Stuttgart
- HEIKE, Georg/THÜRMANN, Eike (1973): *Phonetik*. In: H. P. ALTHAUS/H. HENNE/H.E. WIEGAND (Hg.): *Lexikon der Germanistischen Linguistik*, Studienausgabe Bd. I. Tübingen, 95-105
- HEMPEL, Carl G. (1974): *Grundzüge der Begriffsbildung in der empirischen Wissenschaft*. Düsseldorf (engl. Original 1952)
- HEMPEL, Carl G./OPPENHEIM, Paul (1948): Studies in the Logic of Explanation. In: *Philosophy of Science* 15, 135-175 (abgedruckt in: C. G. HEMPEL: *Aspects of Scientific Explanation*. New York/London 1965, 3-51)
- HENRICI, Gert (1975): Die Binarismus-Problematik in der neueren Linguistik. Tübingen
- HERDAN, Gustav (1969): The mathematical theory of verse. In: *Zeitschrift für Phonetik, Sprachwissenschaft und Kommunikationsforschung* 22, 225-234
- HERDAN, Gustav (1966): *The Advanced Theory of Language as Choice and Chance*. Berlin/Heidelberg/New York
- HERDAN, Gustav (1962): *The Calculus of Linguistic Observations*. The Hague
- HERDAN, Gustav (1954): Informationstheoretische Analyse als Werkzeug der Sprachforschung. In: *Die Naturwissenschaften* 41, 293-295
- HERZOG, Georg (1946): Some linguistic aspects of American Indian poetry. In: *Word* 2, 82
- HEUSLER, Andreas (1956): *Deutsche Versgeschichte*. 3 Bde. Berlin (1925-1929)

- HIRSCHENAUER, Rupert/WEBER, Albrecht (1968): Wege zum Gedicht. Bd. II, Interpretationen von Balladen. München/Zürich
- HOCKETT, Charles F. (1960): A Course in Modern Linguistics. New York 1958
- HÖRMANN, Hans (1970): Psychologie der Sprache. Berlin/Heidelberg/New York
- HOLZKAMP, Klaus (1973): Sinnliche Erkenntnis. Historischer Ursprung und gesellschaftliche Funktion der Wahrnehmung. Frankfurt
- HOLZKAMP, Klaus (1972): Kritische Psychologie. Frankfurt
- HOLZKAMP, Klaus (1968): Wissenschaft als Handlung. Berlin
- HRABÁK, Josef (1966): The Limits of Mathematical Methods in Analyzing Verse. In: *Teorie Verše I* (1966), 265-269
- HRABÁK, Josef (1961): Remarques sur les corrélations entre le vers et la prose, surtout sur les soit-disant formes de transition. In: *Poetics, Poetyka, Poëtika I* (1961), 239-248
- HŘEBÍČEK, Luděk (1965): Euphony in Abay Kunanbayev's Poetry. In: *Asian and African Studies* 1, 123-130
- HRUSHOVSKI, Benjamin (1964): On Free Rhythms in Modern Poetry. In: *SEBEOK* (1964), 173-190
- HUBBELL, Harry M. (1962): Cicero: Orator, Hrsg. von H.M. HUBBELL. London
- HUBER, Egon (1963): Ein Beitrag zur Frage der alternierenden Versbetonung im Französischen. In: *Zeitschrift für französische Sprache und Literatur* 73, 51-58
- IHWE, Jens (1976): Sprache - Struktur - Text - Literaturwissenschaft. In: *ARNOLD/SINEMUS* (1976), 30-44
- IHWE, Jens (1975): On the Foundations of 'Generative Metrics'. In: *Poetics* 16, 367-400
- IHWE, Jens (1972): Linguistik in der Literaturwissenschaft. Zur Entwicklung einer modernen Theorie der Literaturwissenschaft. München
- IHWE, Jens (1971/72): Literaturwissenschaft und Linguistik. 3 Bde. Hrsg. von J. IHWE. Frankfurt
- JACOB, Cary F. (1918): Foundations and Nature of Verse. New York
- JACOBS, Konrad (1970): Turing-Maschinen und zufällige 0-1-Folgen. In: K. JACOBS (Hg.): *Selecta Mathematica II*. Berlin/Heidelberg/New York, 141-167

- JACOBS, Konrad (1969): Maschinenerzeugte 0-1-Folgen. In: K. JACOBS (Hg.): *Selecta Mathematica I*. Berlin/Heidelberg/New York, 1-27
- JÄGER, Siegfried (1972): Linguistik und Statistik. Hrsg. von S. JÄGER. Braunschweig
- JAKOBSON, Roman (1971): Unterbewußte sprachliche Gestaltung in der Dichtung. In: *Zeitschrift für Literaturwissenschaft und Linguistik* 1, 101-112
- JAKOBSON, Roman (1970): The Modular Design of Chinese Regulated Verse. In: *Echanges et Communications. Mélanges offerts à Claude Lévi-Strauss*. The Hague/Paris, 597-605
- JAKOBSON, Roman (1969): Kindersprache, Aphasie und allgemeine Lautgesetze. Frankfurt
- JAKOBSON, Roman (1969a): *Essais de linguistique générale*. Traduit et préfacé par Nicolas Ruwet. Paris 1963
- JAKOBSON, Roman (1966): Über den Versbau der serbokroatischen Volks-epen. In: R. JAKOBSON: *Selected Writings IV*. The Hague, 51-60 (=JAKOBSON 1933)
- JAKOBSON, Roman (1966a): Slavic Epic Verse: Studies in Comparative Metrics. In: R. JAKOBSON: *Selected Writings IV*. The Hague, 414-463
- JAKOBSON, Roman (1964): Linguistics and Poetics. In: *SEBEOK* (1964), 350-377
- JAKOBSON, Roman (1933): Über den Versbau der serbokroatischen Volks-epen. In: *Archives Néerlandaises de Phonétique Expérimentale* 8/9, 44-53 (=JAKOBSON 1966)
- JAKOBSON, Roman (1923): *O českom stixu preimuščestvenno v sopostavlenii s russkim*. Berlin/Moskau
- JAKOBSON, R./FANT, C./HALLE, M. (1951): Preliminaries to Speech Analysis. Cambridge, Mass.
- JAKOBSON, R./HALLE, M. (1970): Phonology in relation to phonetics. In: *MALMBERG* (1970): 411-449
- JAKOBSON, R./LOTZ, J. (1941): Axiome eines Versifikationssystems am mordwinischen Volkslied dargelegt. (Vortrag Stockholm 1941) Englisch Original: Axioms of a versification system. Exemplified by the Mordvinian Folksong. In: *Acta Instituti Hungarici Universitatis Holmiensis, Series B, Linguistica* 1, Stockholm 1952, 5-13 (deutsche Übersetzung in: IHWE (1971/72), Bd. 3, 78-85)
- JAUSS, Hans Robert (1959): Untersuchungen zur mittelalterlichen Tierdichtung. Tübingen

- JOHNSON, Norman/KOTZ, Samuel (1970): Distributions in Statistics: Continuous univariate distributions-1. Boston
- JONES, Daniel (1962): The Phoneme: Its Nature and Use. Cambridge 1950
- JONES, Lawrence Gaylor (1969): Tonality Structure in Russian Verse. In: DOLEŽEL/BAILEY (1969), 122-143
- KAINZ, Friedrich (1967): Psychologie der Sprache. Bd. I. Stuttgart 1941
- KALINKA, Ernst (1937): Griechisch-römische Metrik und Rhythmik im letzten Vierteljahrhundert. In: Bursians Jahresberichte, Supplement Bd. 256, 1-126, Bd. 257, 1-160
- KALINKA, Ernst (1935): Griechisch-römische Metrik und Rhythmik im letzten Vierteljahrhundert. In: Bursians Jahresberichte, Supplement Bd. 250, 290-494
- KANNGIESSER, Siegfried (1976): Kommentar zu ter Meulen. In: D. WUNDERLICH (1976), 101-105
- KAPUR, J. N./SAXENA, H.C. (1970): Mathematical Statistics. New Delhi
- KASSEL, Rudolfus (1965): Aristotelis de arte poetica liber. Ed. R. KASSEL. OXFORD
- KAUFMANN, A. (1975): Introduction to the Theory of Fuzzy Subsets. New York/San Francisco/London
- KAYSER, Wolfgang (1967): Das sprachliche Kunstwerk. Bern 1948
- KENDALL, Maurice G. (1962): Rank Correlation Methods. London 1948
- KENDALL, Maurice G./STUART, Alan (1969): The Advanced Theory of Statistics. London 1969 (1958)
- KERLINGER, Fred N. (1973): Foundations of Behavioral Research. London
- KING, R. D. (1966): On preferred phonemizations for statistical studies: Phoneme frequencies in German. In: Phonetica 15, 22-31
- KINTGEN, E.R. (1974): Is transformational stylistics useful? In: College English 35, 799-824
- KIPARSKY, Paul (1972): Metrics and morphophonemics in the Rigveda. In: M. K. BRAME (Hg.): Contributions to Generative Phonology. Austin, 171-200
- KIPARSKY, Paul (1968): Metrics and morphophonemics in the Kalevala. In: Studies presented to Roman Jakobson by his students. Cambridge, Mass. 1968, 137-148 (abgedruckt in: FREEMAN (1970), 165-181)
- KLEIN, Hans Wilhelm (1966): Phonetik und Phonologie des heutigen Französisch. München 1963

- KLEIN, Wolfgang (1974): Critical Remarks on Generative Metrics. In: Poetics 12, 29-48
- KLEIN, Wolfgang (1974a): Variation in der Sprache. Ein Verfahren zu ihrer Beschreibung. Kronberg/Ts.
- KLIEMANN, Wolfgang/MÜLLER, Norbert (1973): Logik und Mathematik für Sozialwissenschaftler. Bd. I. München
- KLOEPFER, Rolf (1975): Poetik und Linguistik. München
- KLOEPFER, Rolf (1971): Vers libre - Freie Dichtung. In: Zeitschrift für Literaturwissenschaft und Linguistik 3, 81-106
- KLOSTER-JENSEN, Martin (1963): Die Silbe in der Phonetik und Phonemik. In: Phonetica 9, 17-38
- KNAUER, Karl (1965): Die Analyse von Feinstrukturen im sprachlichen Zeitkunstwerk. In: Mathematik und Dichtung (1965), 193-210
- KOCH, Walter A. (1973): Perspektiven der Linguistik I. Hrsg. von W.A. KOCH. Stuttgart
- KOCH, Walter A. (1971): Poetry und Poetizität. In: W. A. KOCH: Varia Semiotica. Hildesheim, 366-395
- KOCH, Walter A. (1966): Recurrence and a three-modal approach to poetry. The Hague
- KOHLER, K.J. (1966): Is the syllable a phonological universal? In: Journal of Linguistics 2, 207-208
- KONDRATOV, A.M. (1969): Information Theory and Poetics: The Entropy of Russian Speech Rhythm. In: DOLEŽEL/BAILEY (1969), 113-121 (zuerst in: Problemy kibernetiki 9 (1963), 279-286)
- KOPCZYŃSKA, Z./MAYENOWA, M.R. (1970): Sur certaines propriétés phonétiques du vers polonais. In: Actes du X^e Congrès international des linguistes, Bd. III, 39-41
- KRÄMER, Wolfgang (1974): Die poetische Destruktion des Phonemischen. In: Kommunikationsforschung und Phonetik. Hamburg, 327-338 (=IPK Forschungsberichte Bd. 50)
- KRÄMER, Wolfgang (1971): Phonetische und phonologische Aspekte des Reims. In: Linguistics 66, 12-28
- KRALLMANN, Dieter (1966): Statistische Methoden in der stilistischen Textanalyse. phil. diss. Bonn
- KRAUSE, Ulrich (1971): Ästhetische Wirkung als Aggregation. In: Zeitschrift für Literaturwissenschaft und Linguistik 4 (1971), 99-114
- KRAUTH, J./LIENERT, G.A. (1973): Die Konfigurationsfrequenzanalyse. Freiburg/München

- KRIZ, Jürgen (1973): Statistik in den Sozialwissenschaften. Reinbek
- KRÖBER, Günter (1968): Der Gesetzesbegriff in der Philosophie und in den Einzelwissenschaften. Hrsg. von G. KRÖBER. Berlin
- KÜPER, Christoph (1973): Möglichkeiten und Grenzen der generativen Metrik. In: Linguistische Berichte 27, 8-40
- KUHN, Thomas S. (1974): Logik der Forschung oder Psychologie der wissenschaftlichen Arbeit? In: I. LAKATOS/A. MUSGRAVE (Hg.): Kritik und Erkenntnisfortschritt. Braunschweig, 1-24
- KUHN, Thomas S. (1970): The Structure of Scientific Revolutions. Chicago
- KULLBACK, S./KUPPERMAN, M./KU, H.H. (1962): An Application of Information Theory to the Analysis of Contingency Tables, With a Table of $2 \times n \ln n$, $n=1(1)10,000$. In: Journal of Research of the National Bureau of Standards - B. Mathematics and Mathematical Physics. Vol. 66B, No. 4, 217-243
- KURYŁOWICZ, Jerzy (1975): Metrik und Sprachgeschichte. Wrocław
- KURYŁOWICZ, Jerzy (1970): Die sprachlichen Grundlagen der altgermanischen Metrik. Innsbruck 1970
- KURYŁOWICZ, Jerzy (1966): A Problem of Germanic Alliteration. In: Studies in language and literature in Honour of M. SCHLAUCH. Warschau, 195-201
- KURYŁOWICZ, Jerzy (1948): Contribution à la théorie de la syllabe. In: Biuletyn Polskiego Towarzystwa Językoznawczego 8, 80-114
- KUTSCHERA, Franz von (1975): Sprachphilosophie. München
- KUTSCHERA, Franz von (1972): Wissenschaftstheorie. 2 Bde. München
- LADEFOGED, Peter (1975): A Course in Phonetics. New York/Chicago
- LAKATOS, Imre (1975): Kritischer Rationalismus und die Methodologie wissenschaftlicher Forschungsprogramme. In: P. WEINGART (Hg.): Wissenschaftsforschung, Frankfurt, 91-132
- LAKATOS, Imre (1974): Falsifikation und Methodologie wissenschaftlicher Forschungsprogramme. In: I. LAKATOS/A. MUSGRAVE (Hg.): Kritik und Erkenntnisfortschritt. Braunschweig, 89-189
- LAUSBERG, Heinrich (1963): Romanische Sprachwissenschaft. Bd. I. Berlin
- LAZICZIUS, Julius von (1935): Probleme der Phonologie. In: Ungarische Jahrbücher 15, 193-208
- LEE, Wayne (1971): Decision Theory and Human Behavior. New York
- LEHFELDT, Werner (1971): Ein Algorithmus zur automatischen Silbentrennung. In: Phonetica 24, 212-237

- LEHISTE, Ilse (1970): Suprasegmentals. Cambridge, Mass.
- LEINFELLNER, Werner (1967): Einführung in die Erkenntnis- und Wissenschaftstheorie. Mannheim/Wien/Zürich
- LENK, Hans (1972): Erklärung, Prognose und Planung. Freiburg
- LÉON, Pierre (1966): Prononciation du français standard. Paris
- LEONT'EV, Aleksej A. (1971): Sprache - Sprechen - Sprechttätigkeit. Stuttgart
- LEPSCHY, Giulio C. (1969): Die strukturelle Sprachwissenschaft. München
- LEUNINGER, H./Miller, M.H./Müller, F. (1972): Psycholinguistik - Ein Forschungsbericht. Frankfurt
- LEVIN, Samuel R. (1971): The Conventions of Poetry. In: S. CHATMAN (Hg.): Literary Style: A Symposium. London/New York, 177-196
- LEVIN, Samuel R. (1963): Deviation - Statistical and Determinate - in Poetic Language. In: Lingua 12, 276-290 (deutsch in: Mathematik und Dichtung (1965), 33-47)
- LEVIN, Samuel R. (1965): Reply to Scholes' "Objections". In: Lingua 13, 193-195
- LEVIN, Samuel R. (1962): Linguistic Structures in Poetry. The Hague/London
- LÉVI-STRAUSS, Claude (1957): Anthropologie Structurale. Paris
- LEVÝ, Jiří (1966): Preliminaries to an Analysis of the Semantic Function of Verse. In: Teorie Verše I (1966), 127-139
- LEVÝ, Jiří (1965): Die Theorie des Verses - ihre mathematischen Aspekte. In: Mathematik und Dichtung (1965), 211-230
- LEVÝ, Jiří (1961): A Contribution to the Typology of Accentual-Syllabic Versification. In: Poetics, Poetyka, Poëtika I (1961), 177-188
- LICHEM, Klaus (1969): Phonetik und Phonologie des heutigen Italienisch. München
- LIGHTFOOT, Marjorie J. (1974): Temporal Prosody: Verse Feet, Measures, Time, Syllabic Distribution, and Isochronous Accent. In: Language and Style 7, 245-260
- LINDSTRAND, Styrbjörn (1943): Os 86 ditongos do Português culto. In: Revista de Portugal 2, Series A, 145-155
- LIST, Gudula (1972): Psycholinguistik. Eine Einführung. Stuttgart
- LIU, James J. Y. (1966): The Art of Chinese Poetry. Chicago
- LOTE, Georges (1919): L'alexandrin d'après la phonétique expérimentale. Paris 1913

- LOTZ, John (1974): Metrics. In: Th. A. SEBEEK (Hg.): Current Trends in Linguistics. Vol. 12. The Hague, 963-982
- LOTZ, John (1972): Elements of Versification. In: WIMSATT (1972), 1-21
- LOTZ, John (1964): Metric Typology. In: SEBEEK (1964), 135-145
- LOTZ, John (1956): A Notation for the Germanic Verse Line. In: Lingua 6, 1-7
- LOTZ, John (1942): Notes on Structural Analysis in Metrics. In: Helicon 4, 119-146
- LÜDTKE, Helmut (1970): Sprache als kybernetisches Phänomen. In: Bibliotheca Phonetica 9, 34-50
- LÜDTKE, Helmut (1968): Metrik und Informationstheorie. In: Theorie Verse II (1968): 61-66
- LÜDTKE, Helmut (1965): Der Vergleich metrischer Schemata hinsichtlich ihrer Redundanz. In: Mathematik und Dichtung (1965), 233-242
- LÜDTKE, Helmut (1953): Fonemática Portuguesa. In: Boletim de Filologia 14, 197-217
- LUELSDOFF, Philip A. (1968): Repetition and Rhyme in Generative Phonology. In: Linguistics 44, 75-90 (deutsche Übersetzung in: IHWE (1971/72), Bd. III, 61-77)
- LYONS, John (1969): Introduction to theoretical linguistics. Cambridge 1968
- MAAS, Heinz Dieter (1971): Einige statistische Untersuchungen zum Werk Georg Trakls. In: Zeitschrift für Literaturwissenschaft und Linguistik 4, 43-50
- MACDOUGALL, R. (1902): The Relation of Auditory Rhythm to Nervous Discharge. In: Psychological Review 9, 460-480
- MAGNUSON, Karl (1974): Rules and Observations in Prosody: Positional Level and Base. In: Poetics 12, 143-154
- MAGNUSON, Karl/Ryder, Frank G. (1971): Second Thoughts on English Prosody. In: College English 33, 198-216
- MAGNUSON, Karl/Ryder, Frank G. (1970): The Study of English Prosody: An Alternative Proposal. In: College English 31, 789-820
- MALMBERG, Bertil (1971): Die expressiven und ästhetischen Ausdrucksmöglichkeiten der Sprache. Ihre strukturelle und quantitative Beschreibung. In: Zeitschrift für Literaturwissenschaft und Linguistik 3, 9-38
- MALMBERG, Bertil (1970): Manual of Phonetics. Hrsg. von B. MALMBERG. Amsterdam 1968

- MALMBERG, Bertil (1967): Structural Linguistics and Human Communication. 1963
- MARCUS, Solomon (1973): Mathematische Poetik. Frankfurt
- MARCUS, Solomon (1967): Entropie et énergie poétique. In: Cahiers de linguistique théorique et appliquée 4, 171-180
- MARKOV, A. A. (1913): Essai d'une recherche statistique sur le texte du roman "Eugène Onegin", illustrant la liaison des épreuves en chaîne. In: Bulletin de l'Académie des Sciences de St. Petersburg 7, 153-162
- MAROUZEAU, Jules (1955): L'e muet dans le vers français. In: Le français moderne 23, 81-86
- MARTIN-LÖF, Per (1966): The Definition of Random Sequences. In: Information and Control 9, 602-616
- MASER, Siegfried (1971): Über das Vermögen einer präzisen Ästhetik. In: Zeitschrift für Literaturwissenschaft und Linguistik 4, 63-72
- Mathematik und Dichtung (1965): Hrsg. von H. KREUZER/R. GUNZENHÄUSER. München
- MAYERHALER, Willi (1974): Einführung in die generative Phonologie. Tübingen
- MAYNTZ, Renate/HOLM, Kurt/HÜBNER, Peter (1972): Einführung in die Methoden der empirischen Soziologie. Opladen 1969
- MCDONALD, B.J. /THOMPSON, W. A. Jr. (1967): Rank sum multiple comparisons in one- and two-way classifications. In: Biometrika 54, 487-497
- MEIER, Helmut (1967): Deutsche Sprachstatistik. Hildesheim
- MELHEM, D.H. (1973): Ivan Fónagy and Paul Delbouille: Sonority Structures in Poetic Language. In: Language and Style 6, 206-215
- MENGES, Günter/SKALA, Heinz J. (1973): Grundriß der Statistik. Teil 2: Daten. Opladen
- MENNEMEIER, Franz Robert (1971): Freier Rhythmus im Ausgang von der Romantik. In: Poetica 4, 197-214
- MEUMANN, Ernst (1894): Untersuchungen zur Psychologie und Ästhetik des Rhythmus. In: Philosophische Studien 10, 249-322, 393-430
- MEUMANN, Ernst (1896): Beiträge zur Psychologie des Zeitbewußtseins. In: Philosophische Studien 12, 127-254
- MEYER-EPPLER, W. (1969): Grundlagen und Anwendungen der Informationstheorie. Berlin 1959
- MILLER, Rupert G. (1966): Simultaneous Statistical Inference. New York
- MOLINO, Jean (1971): La connotation. In: La linguistique 7, 5-30

- MONDOR, H./JEAN-AUBRY, G. (1965): *Oeuvres complètes de Stéphane Mallarmé*. Hrsg. von H. MONDOR/G. JEAN-AUBRY. Paris
- MOOD, Alexander M. (1940): The Distribution Theory of Runs. In: *Annals of Mathematical Statistics* 11, 367-392
- MORAN, P.A.P. (1968): *An Introduction to Probability Theory*. Oxford
- MORCINIEC, N. (1958): Zur phonologischen Wertung der deutschen Affrikaten und Diphthonge. In: *Zeitschrift für Phonetik* 11, 49-66
- MOREAU, R. (1966): Une méthode de décomposition syllabique automatique. In: *Études de Linguistique Appliquée* 4, 65-78
- MORIER, H. (1959): *La psychologie des styles*. Genf
- MORRIS, Charles (1955): *Signs, language, and behavior*. Englewood Cliffs, N.J. 1946
- MÜLLER, Werner (1972): Textklassifikation und Stilanalyse. In: H. SCHANZE (Hg.): *Literatur und Datenverarbeitung*. Tübingen, 160-187
- MULLER, Charles (1964): La longueur moyenne du mot dans le théâtre classique. In: *Cahiers de Lexicologie* 5, 29-44
- MUKAŘOVSKÝ, Jan (1964): Standard Language and Poetic Language. In: GARVIN (1964): 17-33
- MUKAŘOVSKÝ, Jan (1931): La phonologie et la poétique. In: *Travaux du Cercle Linguistique de Prague* 4, 278-288
- NEWMAN, Edwin B. (1951): The Pattern of Vowels and Consonants in Various Languages. In: *American Journal of Psychology* 64, 369-379
- NEWMAN, L./POPPER, W. (1918): *Studies in Biblical Parallelism*. University of California
- NIKITOPOULOS, Pantelis (1972): Qualität, Quantität und Meßbarkeit. In: JÄGER (1972), 105-114
- NILSSON, Nils-Ola (1952): Metrische Stildifferenzen in den Satiren des Horaz. Uppsala
- NÖTH, Winfried (1976): Genese und Arbitrarität der Zeichentypen. In: *Linguistische Berichte* 43, 43-54
- O'CONNOR, J./TRIM, J.L.M. (1953): Vowel, Consonant, and Syllable - a Phonological Definition. In: *Word* 9, 103-122
- OHMANN, Richard (1964): Generative Grammars and the Concept of Literary Style. In: *Word* 20, 423-439 (dt. Übersetzung in: IHWE (1971/72), Bd. I, 213-233)

- OHNESORG, Karel (1966): Le vers enfantin. In: *Teorie Verše I* (1966), 55-58
- OKSAAR, Els (1971): Zum Erwerb des estnischen Phonemsystems in estnisch- und schwedischsprachiger Umgebung. In: *Proceedings of the 9th International Congress of Phonetic Sciences*. Montreal, 746-749
- ONICESCU, Octav (1966): Energie informationnelle. In: *Comptes Rendus de l'Académie des Sciences de Paris* 263, 841-842
- OOMEN, Ursula (1973): *Linguistische Grundlagen poetischer Texte*. Tübingen
- OPP, Karl Dieter (1976): *Methodologie der Sozialwissenschaften. Einführung in Probleme ihrer Theoriebildung*. Reinbek
- OSGOOD, Charles E. (1952): The nature and measurement of meaning. In: *Psychological Bulletin* 49, 197-237
- OSGOOD, C. E./SUCI, C. E./TANNENBAUM, P. H. (1957): *The measurement of meaning*. Urbana
- OTT, Wilhelm (1974): *Metrische Analysen zu Lukrez De Rerum Natura Buch I*. Tübingen
- OTT, Wilhelm (1973): Metrical Analysis of Latin Hexameter - The Automation of a philological Research Project. In: *Linguistica Matematica e Calcolatori*. Firenze, 379-390
- OTT, Wilhelm (1973a): *Metrische Analysen zu Vergil Aeneis Buch I*. Tübingen
- OTT, Wilhelm (1970): *Metrische Analysen zur Ars Poetica des Horaz*. Göttingen
- OWEN, D. B. (1962): *Handbook of Statistical Tables*. Reading, Mass.
- PACE, George B. (1961): The Two Domains: Meter and Rhythm. In: *Publications of the Modern Language Association* 76, 413-419
- PARENT, Monique (1967): *Le vers français au 20^e siècle*. Hrsg. von M. PARENT. Paris
- PASTERNAK, Gerhard (1975): *Theoriebildung in der Literaturwissenschaft*. München
- PAUL, Lothar (1976): Formalisierte Verfahren der Textbeschreibung. In: ARNOLD/SINEMUS (1976), 61-72
- PAUL, Otto/GLIER, Ingeborg (1964): *Deutsche Metrik*. München
- PEARSON, E.S./HARTLEY, H.O. (1970): *Biometrika Tables for Statisticians*. Vol. I. Cambridge 1966

- PETERFALVI, Jean-Michel (1970): Recherches expérimentales sur le symbolisme phonétique. Paris
- PFANZAGL, Johann (1968): Theory of Measurement. Würzburg/Wien
- PHILIPP, Marthe (1968): Phonologie des Graphies et des Rimes. Recherches structurales sur l'alsacien de Thomas MURNER (XVI^e siècle). Paris
- PIEPER, Ursula (1975): Differenzierung von Texten nach numerischen Kriterien. In: Folia Linguistica 7, 61-113
- PIKE, Kenneth L. (1967): Language in Relation to a Unified Theory of the Structure of Human Behavior. The Hague 1959
- PILCH, Herbert (1968): Phonemtheorie. I. Teil. Basel
- Poetics, Poetyka, Poëtika, Bd. II (1966): Hrsg. von R. JAKOBSON u.a. Den Haag
- Poetics, Poetyka, Poëtika, Bd. I (1961): Hrsg. von D. DAVIE u.a. Warschau/Den Haag
- POPPE, Nikolaus (1958): Der Parallelismus in der epischen Dichtung der Mongolen. In: Ural-Altaische Jahrbücher 30, 195-228
- POPPER, Karl R. (1974): The Poverty of Historicism. London 1957
- POPPER, Karl R. (1973): Logik der Forschung. Tübingen 1934
- POPPER, Karl R. (1972): Truth, Rationality and the Growth of Knowledge. In: POPPER (1972b), 215-250
- POPPER, Karl R. (1972a): Objective Knowledge. An Evolutionary Approach. Oxford (dt. Übersetzung: Frankfurt 1973)
- POPPER, Karl R. (1972b): Conjectures and Refutations. The Growth of Scientific Knowledge. London 1963
- POPPER, Karl R. (1957): Über die Zielsetzung der Erfahrungswissenschaft. In: Ratio 1, 21-31 (abgedruckt in: H. ALBERT (Hg.): Theorie und Realität. Tübingen 1972, 29-41)
- PULGRAM, Ernst (1970): Syllable, Word, Nexus, Cursus. The Hague
- RADNITZKY, Gerard (1973): Contemporary Schools of Metascience. Chicago
- RAUHUT, F. (1935): Die Entstehung des Prinzips der Silbenzahl in der französischen Verskunst. In: Germanisch-romanische Monatsschrift 13, 129-146
- REGNAUD, R. (1884): La rhétorique sanscrite. Paris
- REICHENBACH, Hans (1938): Experience and Prediction. An Analysis of the Foundations and the Structure of Knowledge. Chicago (61966)

- REISS, Katharina (1971): Möglichkeiten und Grenzen der Übersetzungskritik. München
- RESCHER, Nicholas (1970): Scientific Explanation. New York
- RIEGER, Burghard (1972): Warum mengenorientierte Textwissenschaft? Zur Begründung der Statistik als Methode. In: Zeitschrift für Literaturwissenschaft und Linguistik 8, 11-28
- RIESER, Hannes/WIRRER, Jan (1974): Zu Teun van Dijk's "Some Aspects of Text Grammars". Ein Beitrag zur Textgrammatik und Literaturdiskussion. In: Probleme und Perspektiven der neueren textgrammatischen Forschung I. Hrsg. von der Projektgruppe Textlinguistik Konstanz. Hamburg, 1-80
- RIFFATERRE, Michael (1971): Essais de stylistique structurale. Paris (dt. Übersetzung: München 1973)
- RIFFATERRE, Michael (1960): Stylistic Context. In: Word 16, 207-218
- ROACH, S.A. (1968): The Theory of Random Clumping. London
- ROUBAUD, Jacques (1971): Mètre et vers. In: Poétique 2, 366-387
- RUDNER, Richard S. (1966): Philosophy of Social Science. Englewood Cliffs, N.J.
- RUIPÉREZ, M. (1955): Cantidad silábica y métrica estructural en griego antiguo. In: Emerita 23, 79-95
- RŮŽIČKA, Rudolf (1970): Die Begriffe 'merkmalhaltig' und 'merkmallos' und ihre Verwendung in der generativen Transformationsgrammatik. In: BIERWISCH/HEIDOLPH (1970), 260-284
- RYDER, Frank G. (1967): Vowels and consonants as features of style. Some poems of Goethe and Klopstock. In: Linguistics 37, 89-110
- ŠABRŠULA, Jan (1969): Le signe linguistique et la poésie versifiée. In: Folia Linguistica 5, 182-184
- SACHS, Lothar (1974): Angewandte Statistik. Berlin
- SALOMAA, Arto (1969): Probabilistic and Weighted Grammar. In: Information and Control 15, 529-544
- SANDERS, Willy (1973): Linguistische Stiltheorie. Göttingen
- SAUSSURE, Ferdinand de (1968): Cours de linguistique générale. Paris 1916
- SCHÄDLICH, Hans-Joachim (1969): Über Phonologie und Poetik. In: Jahrbuch für Internationale Germanistik 1, 44-60 (abgedruckt in: IHWE (1971/72), Bd. III, 42-60)
- SCHEEL, Hans Ludwig (1970): Zur Theorie und Praxis der Versforschung II. In: Romanistisches Jahrbuch 21, 54-74

- CHEEL, Hans Ludwig (1967): Zur Theorie und Praxis der Versforschung I. In: Romanistisches Jahrbuch 18, 38-55
- CHEUCH, Erwin K./ZEHNPFENNIG, Helmut (1974): Skalierungsverfahren in der Sozialforschung. In: R. KÖNIG (Hg.): Handbuch der empirischen Sozialforschung. Bd. 3a. Stuttgart, 97-203
- CHMIDT, Julius (1968): Der Rhythmus des französischen Verses. Wiesbaden (= Zeitschrift für französische Sprache und Literatur, Beiheft 2)
- CHMIDT, Siegfried J. (1975): Literaturwissenschaft als argumentierende Wissenschaft. München
- CHMIDT, Siegfried J. (1973): Texttheorie. Probleme einer Linguistik der sprachlichen Kommunikation. München
- CHMIDT, Siegfried J. (1973a): Texttheorie/Pragmalinguistik. In: H.P. ALTHAUS/H. HENNE/H.E. WIEGAND (Hg.): Lexikon der Germanistischen Linguistik, Studienausgabe Bd. II. Tübingen, 233-244
- CHNORR, Claus Peter (1971): Zufälligkeit und Wahrscheinlichkeit. Berlin/Heidelberg/New York
- CHRÖDER, Hartwig (1960): Quantitative Stilanalyse. Versuch einer Analyse quantitativer Stilmerkmale unter psychologischem Aspekt. Diss. Würzburg
- CHULTZ, Hartwig (1972): Methoden und Aufgaben einer zukünftigen Metrik. In: Sprache im technischen Zeitalter 41, 27-51
- CHWARTZ, E./WIMSATT, W.K./BEARDSLEY, M.C. (1962): Rhythm and 'Exercises in Abstraction'. In: Publications of the Modern Language Association 77, 668-674
- CRIPTURE, E.W. (1929): Grundzüge der englischen Verswissenschaft. Marburg
- EASHORE, Carl E./METFESSEL, Milton (1925): Deviation from the Regular as an Art Principle. In: Proceedings of the National Academy of Sciences of the United States of America 11, 538-542
- SEBEOK, Thomas A. (1964): Style in Language. Hrsg. von Th. A. SEBEOK. Cambridge, Mass. 1960
- SEBEOK, Thomas A. (1963): The Informational Model of Language: Analog and Digital Coding in Animal and Human Communication. In: P. GARVIN (Hg.): Natural Language and the Computer. New York 1963, 47-64
- SEBEOK, Th. A./ZEPS, V.J. (1959): On Non-Random Distribution of Initial Phonemes in Cheremis Verse. In: Lingua 8, 370-384

- SEREBRJANNIKOV, D.F. (1974): Heuristische Prinzipien und logische Kalküle. München/Salzburg (russ. Original 1970)
- SERVIEN, Pius (1930): Les rythmes comme introduction physique à l'esthétique. Paris
- SHANNON, Claude E./WEAVER, Warren (1949): The Mathematical Theory of Communication. Urbana
- SHAPIRO, Karl (1948): A Bibliography of Modern Prosody. Baltimore
- SIEVEKE, Franz Günter (1973): Metrik als Theorie phonomorpher Wirkmuster. In: D. BREUER u.a. (Hg.): Literaturwissenschaft. Eine Einführung für Germanisten. Frankfurt/Berlin/Wien, 341-390
- SIEVERS, Eduard (1893): Grundzüge der Phonetik. Leipzig
- SKINNER, B.F. (1941): A Quantitative Estimate of Certain Types of Sound-Patterning in Poetry. In: The American Journal of Psychology 54, 64-79
- SKINNER, B.F. (1939): The Alliteration in Shakespeare's Sonnets. A Study in Literary Behavior. In: The Psychological Record 3, 186-192
- SLAMA-CAZACU, Tatiana (1967): Sur les rapports entre la stylistique et la psycholinguistique. In: Revue roumaine de linguistique 12, 309-330
- SNIDER, James G./OSGOOD, Charles E. (1969): Semantic Differential Technique: A Sourcebook. Hrsg. von J.G. SNIDER/C.E. OSGOOD, Chicago
- SOMERS, H.H. (1967): Analyse statistique du style. Louvain/Paris
- SONNENSCHNITT, Edward (1925): What is Rhythm? Oxford
- SPINNER, Helmut (1974): Pluralismus als Erkenntnismodell. Frankfurt
- STACHOWIAK, Herbert (1973): Allgemeine Modelltheorie. Wien/New York
- STACHOWIAK, Herbert (1965): Gedanken zu einer allgemeinen Theorie der Modelle. In: Studium Generale 18, 432-463
- STANDOP, Ewald (1972): Die Metrik auf Abwegen. Eine Kritik der Halle-Keyser-Theorie. In: Linguistische Berichte 19, 1-19
- STANKIEWICZ, Edward (1974): Structural Poetics and Linguistics. In: Th. A. SEBEOK (Hg.): Current Trends in Linguistics. Vol.12. The Hague, 629-659
- STANKIEWICZ, Edward (1964): Linguistics and the Study of Poetic Language. In: SEBEOK (1964), 69-81
- STANLEY, George Edward (1971): Phonoaesthetics and West Texas Dialect. In: Linguistics 71, 95-102

- STECHOW, Arnim von (1970): Aspekte zur Bewertung von generativen Grammatiken. In: Linguistische Berichte 9, 18-28
- STEFENELLI-FÜRST, Friederike (1966): Die Tempora der Vergangenheit in der Chanson de Geste. Wien
- STEGMÜLLER, Wolfgang (1969-1974): Probleme und Resultate der Wissenschaftstheorie und Analytischen Philosophie. 4 Bde. Berlin/Heidelberg/New York
- STEGMÜLLER, Wolfgang (1974): Theorie und Erfahrung. Erster Halbband: Begriffsformen, Wissenschaftssprache, empirische Signifikanz und theoretische Begriffe. Berlin (=STEGMÜLLER 1969-1974, Bd. II)
- STEGMÜLLER, Wolfgang (1973): Theorie und Erfahrung. Zweiter Halbband: Theorienstrukturen und Theoriendynamik. Berlin (=STEGMÜLLER 1969-1974, Bd. II)
- STEGMÜLLER, Wolfgang (1973a): Personelle und Statistische Wahrscheinlichkeit. Zweiter Halbband: Statistisches Schließen, Statistische Begründung, Statistische Analyse. Berlin (=STEGMÜLLER 1969-1974, Bd. IV)
- STEGMÜLLER, Wolfgang (1969): Wissenschaftliche Erklärung und Begründung. Berlin (=STEGMÜLLER 1969-1974, Bd. I)
- STEGMÜLLER, Wolfgang (1965): Hauptströmungen der Gegenwartsphilosophie. Stuttgart
- STEINITZ, Wolfgang (1934): Der Parallelismus in der finnisch-karelischen Volksdichtung. Helsinki (=Folklore Fellows Communications 115)
- STEINTHAL, H. (1961): Geschichte der Sprachwissenschaft bei den Griechen und Römern. Bd. I. Berlin ²1890
- STETSON, R.H. (1903): Rhythm and Rhyme. In: The Psychological Review, Monograph Supplements 4, 413-466
- STETSON, R.H. (1928): Motor Phonetics. In: Archives néerlandaises de Phonétique expérimentale 3, 1-216
- STEVENS, Stanley S. (1946): On the Theory of Scales of Measurement. In: Science 103, 677-680
- STEVENS, Stanley S. (1959): Measurement, psychophysics and utility. In: C.W. CHURCHMAN/Ph. RATOOSH (Hg.): Measurement, Definitions and Theories. New York, 18-64
- ŠTOKMAR, M. P. (1933): Bibliografija rabot po stixosloženiju. Moskau
- ŠTUKOVSKÝ, Robert/ALTMANN, Gabriel (1966): Die Entwicklung des slowakischen Reims im XIX Jahrhundert. In: Theorie Verše I (1966), 259-261
- SUPPES, Patrick (1972): Probabilistic Grammars for Natural Languages. In: D. DAVIDSON/G. HARMAN (Hg.): Semantics of Natural Language. Dordrecht, 741-762

- SUPPES, Patrick/ZINNES, Joseph L. (1963): Basic Measurement Theory. In: R.D. LUCE/R.R. BUSH/E. GALANTER (Hg.): Handbook of Mathematical Psychology. Bd. I. New York 1963, 1-76
- SUTTERHEIM, Cornelius F.P. (1961): Poetry and Prose, their Interrelations and Transitional Forms. In: Poetics, Poetyka, Poëtika I (1961), 225-237
- SWED, Frieda S./EISENHART, C. (1943): Tables for testing randomness of grouping in a sequence of alternatives. In: Annals of Mathematical Statistics 14, 66-87
- TARANOVSKI, Kiril (1963): Metrics. In: Th. A. SEBEEK (Hg.): Current Trends in Linguistics. Vol I, Soviet and East European Linguistics. Den Haag, 192-201
- TARLINSKAJA, M.G. (1974): Meter and Rhythm of Pre-Chaucerian Rhymed Verse. In: Linguistics 121, 65-87
- TARLINSKAJA, M.G./TETERINA, L.M. (1974): Verse - Prose - Metre. In: Linguistics 129, 63-86
- TARNÓCZY, Thomas (1961): Phonetische Gesichtspunkte bei der Zusammenstellung von Texten für Verständlichkeitsmessungen. In: Zeitschrift für Phonetik 14, 74-87
- TARSKI, Alfred (1935): Der Wahrheitsbegriff in den formalisierten Sprachen. In: Studia Philosophica Comentarior Societas Philosophicae Polonorum 1, 261-405 (abgedruckt in: K. BERKA/L. KREISER (Hg.): Logiktexte. Berlin 1971, 447-559)
- TARSKI, Alfred (1944): The Semantic Conception of Truth and the Foundations of Semantics. In: Philosophy and Phenomenological Research 4, 341-376 (deutsche Übersetzung in: J. SINNREICH (Hg.): Zur Philosophie der idealen Sprache. München 1972, 53-100)
- TAYLOR, W.L. (1953): Cloze procedure, a new tool for measuring readability. In: Journalism Quarterly 30, 415-433
- Teorie Verše II (1968) Hrsg. von J. LEVÝ/K. PALAS. Brno
- Teorie Verše I (1966): Hrsg. von J. LEVÝ. Brno
- TER MEULEN, Alice (1976): Grammars and empirical theories. In: WUNDERLICH (1976), 87-100
- THOMPSON, John (1961): Linguistic structure and the poetic line. In: Poetics, Poetyka, Poëtika I (1961), 167-175 (abgedruckt in: FREEMAN (1970), 336-346)
- THIEME, Hugo (1933): Bibliographie de la littérature française. De 1800 à 1930. Paris
- THORNE, James P. (1969): Poetry, Stylistics and Imaginary Grammars. In: Journal of Linguistics 5, 147-150

- THORNE, James P. (1965): Stylistics and Generative Grammar. In: *Journal of Linguistics* 1, 49-59
- TILLMANN, Hans Günter (1964): Das phonetische Silbenproblem. Eine theoretische Untersuchung. Diss. Bonn
- TORGERSON, Warren S. (1958): Theories and Methods of Scaling. New York/London
- TROJAN, Felix (1975): Biophonetik. Hrsg. von H. SCHENDL. Mannheim/Wien/Zürich
- TRUBETZKOY, N.S. (1967): Grundzüge der Phonologie. Göttingen ¹1958
- TRUNZ, Erich (1962): Goethes Werke. Bd. 1. Hamburg ¹1948
- TURČÁNY, Viliam (1966): Zamečania po voprosu vzaimosvjazej različnych elementon poetičeskobo proizvedenija. (Bemerkungen zur Frage der Wechselbeziehung zwischen den einzelnen Elementen eines poetischen Kunstwerks.) In: *Teorie Verše* I (1966), 143-150
- TURNER, Ronald C. (1974): An Automated Procedure for Quantification of Rhythmical Patterning in Spanish. In: *Linguistics* 121, 89-98
- UNGEHEUER, Gerold (1969): Das Phonemsystem der deutschen Hochlautung. In: H. de BOOR/H. MOSER/Ch. WINKLER (Hg.): *Siebs Deutsche Aussprache*. Berlin, 27-42
- UNGEHEUER, Gerold (1965): Extensional-paradigmatische Bestimmung auditiver Qualitäten phonetischer Signale. In: *Proceedings of the Fifth International Congress of Phonetic Sciences*. Basel, 556-560
- VALESIO, Paolo (1971): On Poetics and Metrical Theory. In: *Poetics* 2, 36-60
- VALESIO, Paolo (1966): Problemi di metrica: un esperimento di tipologia italo-russa. In: *Lingua e Stile* 1, 305-321
- VAN TIEGHEM, Philippe (1965): Les grandes doctrines littéraires en France. Paris
- VENNEMANN GENANT NIERFELD, Theo (1972): On the Theory of Syllabic Phonology. In: *Linguistische Berichte* 18, 1-18
- VERRIER, Paul (1931-1932): *Le vers français*. 3 Bde. Paris
- VERRIER, Paul (1909): *Essai sur les principes de la métrique anglaise*. 3 Bde. Paris
- VERWEY, Albert (1934): *Rhythmus und Metrum*. Halle
- VOEGELIN, C.F./EULER, R.C. (1957): Introduction to Hopi Chants. In: *Journal of American Folklore* 70, 115-136

- WALLIN, Wallace J.E. (1901): Researchs on the Rhythm of Speech. In: *Studies from the Yale Psychological Laboratory* 9, 1-142
- WALLIS, W. Allen (1952): Rough-and-Ready Statistical Tests. In: *Industrial Quality Control* 8, 35-40
- WANG, Jün-tin (1972): Wissenschaftliche Erklärung und generative Grammatik. In: K. HYLDGAARD-JENSEN (Hg.): *Linguistik 1971: Referate des 6. Linguistischen Kolloquiums*, 11. - 14.8.1971, Kopenhagen. Frankfurt, 50-66
- WEAVER, Warren (1948): Probability, Rarity, Interest, and Surprise. In: *The Scientific Monthly* 67, 390-392
- WELLEK, René/WARREN, Austin (1966): *Theorie der Literatur*. Berlin
- WELLS, Rulon (1964): Comments to Part Five: Metrics. In: SEBEOK (1964), 197-200
- WESTPHAL, R. (1965): Aristoxenos von Tarent. Melik und Rhythmik des Classischen Hellenentums. Bd. 2. Hildesheim ¹1893
- WICKMANN, Dieter (1972): Urteilsstruktur und Signifikanzschwelle. Quantifizierung eines nicht-numerischen Problems: Statistik zur unbekannten Verfasserschaft. In: H. SCHANZE (Hg.): *Literatur und Datenverarbeitung*. Tübingen, 107-122
- WIENOLD, Götz (1972): *Semiotik der Literatur*. Frankfurt
- WIENOLD, Götz (1971): Formulierungstheorie - Poetik - Strukturelle Literaturgeschichte. Frankfurt
- WILCOXON, F./WILCOX, R.A. (1964): *Some Rapid Approximate Statistical Procedures*. New York
- WILKINS, A.S. (1963): *M. Tulli Ciceronis Rhetorica. Tomus I. ed. A.S. Wilkins*. Oxford
- WILKS, Samuel S. (1962): *Mathematical Statistics*. New York
- WILLIAMS, C.B. (1970): *Style and Vocabulary*. London
- WIMSATT, William K. (1972): Versification. *Modern Language Types. Sixteen Essays*. Hrsg. von W.K. WIMSATT. New York
- WIMSATT, W.K. Jr./BEARDSLEY, M.C. (1964): The Concept of Meter: an Exercise in Abstraction. In: SEBEOK (1964), 193-196
- WIMSATT, W.K. Jr./BEARDSLEY, M.C. (1959): The Concept of Meter: an Exercise in Abstraction. In: *Publications of the Modern Language Association of America* 74, 585-598
- WINCKEL, Fritz (1960): *Phänomens des musikalischen Hörens*. Berlin
- WINTERS, Yvor (1962): The function of criticism. Problems and exercises. London ¹1957

- WITTSACK, Walter (1962): Zur Frage der Klangform des Naturmagischen in Goethes "Erlkönig". In: Wissenschaftliche Zeitschrift der Martin-Luther-Universität Halle-Wittenberg. Gesellschafts- und sprachwissenschaftliche Reihe. Jg. XI, 12, 1773-1778
- WODE, Henning (1971): Linguistische Grundlagen verslicher Strukturen im Englischen. In: Folia Linguistica 4, 372-392
- WOLF, Willi (1974): Statistik I. Eine Einführung für Sozialwissenschaftler. Weinheim/Basel
- WOODROW, Herbert (1951): Time Perception. In: S.S. STEVENS (Hg.): Handbook of Experimental Psychology. New York/London, 1224-1236
- WORONCZAK, Jerzy (1961): Statistische Methoden in der Verslehre. In: Poetics, Poetyka, Poëtika (1961), 607-624
- WRIGHT, Georg Henrik von (1974): Erklären und Verstehen. Frankfurt (engl. Original 1971)
- WUNDERLICH, Dieter (1976): Wissenschaftstheorie der Linguistik. Hrg. von D. WUNDERLICH. Kronberg
- WUNDERLICH, Dieter (1974): Grundlagen der Linguistik. Reinbek
- WUNDERLICH, Dieter (1970): Die Rolle der Pragmatik in der Linguistik. In: Der Deutschunterricht 22, 5-41
- WUNDT, Wilhelm (1900): Völkerpsychologie, Bd. I: Die Sprache. Leipzig
- YODER, Perry B. (1972): Biblical Hebrew. In: WIMSATT (1972), 52-65
- ZIPF, George Kingsley (1965): The Psycho-Biology of Language. Cambridge. Mass. 1935
- ZIRIN, Ronald A. (1970): The Phonological Basis of Latin Prosody. The Hague
- ŽIRMUNSKIJ, Viktor (1966): Introduction to Metrics. The Hague (russ. Original 1925)