

WOJCIECH SKALMOWSKI

ÜBER EINIGE STATISTISCH ERFASSBARE ZÜGE DER PERSISCHEN SPRACHENTWICKLUNG *

Herrn Prof. Dr. Heinrich F. J. Junker
in Dankbarkeit und Verehrung gewidmet

1. THEORETISCHE GRUNDLAGEN DER UNTERSUCHUNG

1.1. Ziel der Arbeit

Die vorliegende Arbeit hat als Ziel die statistischen Verteilungen der Wörter von unterschiedlicher Anzahl der Silben und Phoneme in den drei Entwicklungsstufen der persischen Sprache, d. h. im Alt-, Mittel- und Neupersischen zu untersuchen. Die gewonnenen Daten dienen als Material bei einem Versuch einer allgemeinen Charakterisierung der persischen Sprachentwicklung, welche vom Standpunkt mancher Erkenntnisse und Forderungen der Informationstheorie durchgeführt wird. Man bedient sich in dieser Untersuchung eines statistischen Modells der sprachlichen

* Der vorliegende Aufsatz bildet eine Umarbeitung der Dissertation des Vf., die u. d. T. *Sprachstatistische Untersuchungen zur persischen Sprachentwicklung* 1960 der Philosophischen Fakultät der Humboldt-Universität zu Berlin vorgelegt und von ihr genehmigt wurde. Die Anregung zu dieser Arbeit sowie auch wertvolle kritische Hinweise verdankt der Vf. dem Direktor des Vorderasiatischen Institutes der Humboldt-Universität, Herrn Prof. Dr. H. F. J. Junker. Ein besonderer Dank gilt auch dem Leiter der Abteilung für allgemeine Sprachwissenschaft der Jagellonischen Universität zu Kraków, Herrn Prof. Dr. Jerzy Kuryłowicz, der diese Dissertation in Manuskript gelesen und begutachtet hat.

Kommunikation, das im grossen und ganzen in Anlehnung an die Sprachauffassung von de Saussure konstruiert wird; den Ausgangspunkt bilden hierbei: 1° seine Auffassung des sprachlichen Zeichens als eines Ganzen, welches das Bezeichnende (signifiant) und das Bezeichnete (signifié) umfasst, und 2° seine Aufteilung der menschlichen Rede (langage) in die Sprache (langue) und das Sprechen (parole).

Die vorliegende Arbeit befasst sich ausschliesslich mit der signifiant-Seite des sprachlichen Zeichens, d. h. sie betrachtet den Sprachcode als ein System der Symbole, die *ex definitione* eine signifié-Seite besitzen.

1.2. Statistisches Modell der sprachlichen Kommunikation

Das sprachliche Zeichen hat zwei Grundeigenschaften, welche de Saussure folgendermassen definiert:

1. „Das sprachliche Zeichen ist beliebig“¹;
2. „Das Bezeichnende (signifiant), als etwas Hörbares, verläuft ausschliesslich in der Zeit und hat Eigenschaften, die von der Zeit bestimmt sind:

- a) es stellt eine Ausdehnung dar, und
- b) diese Ausdehnung ist messbar in einer einzigen Dimension; es ist eine Linie“².

Die erste Grundeigenschaft weist der Sprache eine Stellung unter den symbolischen Coden an, welche aus diskreten Einheiten bestehen. Diese Einheiten werden Symbole genannt. Die bedeutungstragenden Symbole können sich aus einer endlichen Anzahl von elementaren Symbolen zusammensetzen. Die Zusammensetzung darf nur in einer Dimension erfolgen; deswegen — bei endlicher Anzahl der elementaren Symbole — können die bedeutungstragenden Symbole nur nach den kombinatorischen Prinzipien aufgebaut werden. Die Verschiedenheit der Symbole beruht

¹ F. de Saussure, *Grundfragen der allgemeinen Sprachwissenschaft*, übers. von H. Lommel, Berlin u. Leipzig 1931, S. 79.

² Ebenda, S. 82.

in diesem Falle auf der Verschiedenheit der Anordnung der elementaren Symbole in einer Reihenfolge.

Die Aneinanderreihung der Symbole unterliegt verschiedenen Einschränkungen, die u. a. durch die physikalischen Bedingungen der Realisierung der Symbole (z. B. der Phoneme), die gegenseitige Beeinflussung des Bezeichneten und des Bezeichnenden u. s. w. verursacht werden. Aus diesem Grunde hat jedes sprachliche System eine bestimmte Struktur, die u. a. in den statistischen Eigenschaften dieses Systems zum Ausdruck kommt.

Eine dieser Eigenschaften ist die relative Häufigkeit des Auftretens der Symbole und ihrer Gruppen. Sie lässt sich zahlenmässig erfassen und ermöglicht dadurch eine Beschreibung des Systems mit Hilfe quantitativer Grössen. Die relativen Häufigkeiten des Auftretens der einzelnen Sprachelemente in einer gegebenen Entwicklungsstufe einer Sprache besitzen einen festen Wert. Diese Eigenschaft der Sprache fand längst eine praktische Verwendung in der Kryptographie. Die Kenntnis der relativen Häufigkeiten der Buchstaben bzw. der Buchstabengruppen ermöglicht die Entzifferung der sogen. Substitutionschiffre¹. Die Chiffren dieser Art beruhen darauf, dass die Buchstaben bzw. Buchstabengruppen eines Textes, der für Uneingeweihte geheim bleiben soll, nach bestimmter Regel durch andere Symbole ersetzt werden. Das Geheimnis einer solchen Schrift liegt in der Unkenntnis, welches Chiffersymbol welchem Buchstaben entspricht. Derartige Chiffertexte lassen sich ohne Schlüssel entziffern — ihre Schwäche liegt darin, dass die Buchstaben zwar ihre allgemeinbekannte Form verlieren, aber ihre relativen Häufigkeiten, nach welchen sie identifiziert werden können, beibehalten. Man kann darin volle Analogie mit dem sprachlichen Zeichen ersehen: die Buchstaben entsprechen dem Bezeichneten, die Chiffersymbole dem Bezeichnenden; in beiden Fällen ist die relative Häufigkeit — in der Terminologie der Wahrscheinlichkeitstheorie: die Wahrscheinlichkeit — dasjenige Merkmal des Zeichens, das seine beiden Seiten miteinander verbindet.

In diesem Zusammenhang ist es interessant den Begriff des sprachlichen Wertes (valeur) zu betrachten, der von de Saussure

¹ Vgl. L. D. Smith, *Cryptography*, N. Y. 1943, S. 97 ff.

sure folgendermassen definiert wird: „(die Werte) sind immer gebildet:

1. durch etwas Unähnliches, das ausgewechselt werden kann gegen dasjenige, dessen Wert zu bestimmen ist;
2. durch ähnliche Dinge, die man vergleichen kann mit demjenigen, dessen Wert in Rede steht“¹.

Der sprachliche Wert ermöglicht den Austausch des Bezeichneten gegen das Bezeichnende und *vice versa* innerhalb des Zeichens und lässt einen das eine Zeichen mit einem anderen vergleichen — er spielt also dieselbe Rolle, wie die Wahrscheinlichkeit des Zeichens im gesamten System. Diese enge Beziehung zwischen den beiden Begriffen gestattet es, den Wert als eine Funktion der Wahrscheinlichkeit zu betrachten. Damit wird der Wert in einen quantitativen Begriff umgewandelt.

De Saussure nennt die Gesamtheit der sprachlichen Kommunikation die menschliche Rede (*langage*) und teilt sie in die Sprache (*langue*) und das Sprechen (*parole*)².

Die Sprache wird folgendermassen definiert: „... es ist die Gesamtheit der sprachlichen Gewohnheiten, welche es dem Individuum gestatten, zu verstehen und sich verständlich zu machen“³.

„Die Sprache besteht in der Sprachgemeinschaft in Gestalt einer Summe von Eindrücken, die in jedem Gehirn niedergelegt sind, ungefähr so wie ein Wörterbuch, von dem alle Exemplare, unter sich völlig gleich, unter den Individuen verteilt wären...“⁴.

Das Sprechen dagegen ist: „die Summe von allem, was die Sprachgenossen reden, und umfasst: a) die individuellen Kombinationen, welche abhängig sind von dem Willen des Sprechenden, b) die Akte der Lautgebung, welche gleichermassen vom Willen bestimmt werden und notwendig sind zur Verwirklichung jener Kombinationen“⁵. „Indem man die Sprache vom Sprechen scheidet, scheidet man zugleich: 1) das Soziale vom Individuellen; 2) das Wesentliche vom Akzessorischen und mehr oder weniger Zufälligen“⁶.

¹ De Saussure, a. a. O., S. 137. ² Ebenda, S. 13ff.

³ Ebenda, S. 91. ⁴ Ebenda, S. 23. ⁵ Ebenda, S. 23.

⁶ Ebenda, S. 16.

Seine Erwägungen über die Unterschiede zwischen den beiden Teilen der menschlichen Rede abschliessend stellt de Saussure fest: „aus allen diesen Gründen ist es nicht möglich, Sprache und Sprechen unter einem und demselben Gesichtspunkt zu vereinigen“¹.

Diese Feststellung ist jedoch unvereinbar mit dem Begriff des sprachlichen Wertes, der einerseits in dem Sprechen und durch das Sprechen realisiert wird, andererseits, als ein untrennbarer Teil des Zeichens, der Sprache angehört; s. dazu de Saussure, a. a. O., S. 135 und S. 150.

Verbindet man das Bild der Sprache als eines Wörterbuches von Zeichen mit der Auffassung des Wertes als einer Funktion der Wahrscheinlichkeit dieser Zeichen, so gelangt man zu einem statistischen Aspekt der menschlichen Rede², deren beide Teile — die Sprache und das Sprechen — sich so zueinander verhalten, wie die Verteilung eines Kollektivs zu den Proben, die diesem Kollektiv aufs Geradewohl entnommen werden.

Der Begriff Kollektiv gehört der Wahrscheinlichkeitstheorie an. Unter ihm werden Erscheinungen oder Vorgänge verstanden, „die zwei Forderungen genügen, nämlich der nach der Existenz eines Grenzwertes und der nach der ‘Regellosigkeit’ der Zuordnung. Sobald von diesen Merkmalen nur zwei vorhanden sind. z. B. ‘0 oder 1’, also das einfachste Kollektiv, bezeichnet man dies als Alternative. Wenn die Häufigkeit des Auftretens einzelner Ziffern bestimmten Grenzwerten zustrebt, bezeichnet man diese auch als ‘Wahrscheinlichkeiten’. Die Gesamtheit der Wahrscheinlichkeitswerte innerhalb eines Kollektivs wird ‘Verteilung’ genannt“³.

Betrachtet man die menschliche Rede als eine Menge regellos verteilter Zeichen, so kann man in ihr ein Kollektiv sehen, dessen Verteilung die Sprache (*langue*) bildet, die man als eine Norm auffassen kann. Man kann sie sich vorstellen als ein „Wörterbuch“ von Zeichen mit ihren Werten versehen; auf den Zusammenhang zwischen dem Wert und der Wahrscheinlichkeit wurde bereits oben

¹ Ebenda, S. 23.

² G. Herdan, *Language as Choice and Chance*, Groningen 1956, S. 79.

³ P. Neidhardt, *Einführung in die Informationstheorie*, Berlin—Stuttgart 1957, S. 15.

hingewiesen. Die in dem Sprechen vorkommenden Abweichungen von der Norm entsprechen den zufälligen Beobachtungsfehlern in den Proben aus einem Kollektiv. Man beachtet, dass die Kriterien, nach denen de Saussure die beiden Teile der menschlichen Rede voneinander trennt, nämlich: 1) das Soziale — das Individuelle, 2) das Wesentliche — das Zufällige, den gegenseitigen Beziehungen zwischen dem Kollektiv und den Proben genau entsprechen. Die Sprache, als Gesamtheit der Wahrscheinlichkeitswerte aller Bestandteile eines Kollektivs aufgefassen, stellt eine Norm dar, die eine Entscheidung möglich macht, ob die dargelegten Proben — d. h. das Sprechen — dem gegebenen Kollektiv angehören. Als eine statistische Norm lässt sie in den Proben individuelle Abweichungen bestimmter Grösse von der mittleren Verteilung zu, ohne dass die Identität der Proben mit dem Kollektiv bestritten wird. Die Grösse der zulässigen Abweichungen ist von der Struktur des Kollektivs abhängig; andererseits im Falle der menschlichen Rede ist die Verteilung des Kollektivs in Wirklichkeit etwas Abstraktes. Es realisiert sich durch das Sprechen und wird immer wieder aufs Neue geschaffen. Auf diese Weise kann sich die neu-entstehende Verteilung von der vorhergehenden unterscheiden, wobei die Kontinuität der Verständigung nicht unterbrochen wird, weil sich die Verschiebung der Verteilung innerhalb der zulässigen Fehlergrenze vollzogen hat. Eine derartige Darstellung der sprachlichen Kommunikation, in welcher das Sprechen und die Sprache in einer gegenseitigen Abhängigkeit voneinander stehen, lässt sie als ein mit Hilfe der Rückkopplung selbstgesteuertes System im kybernetischen Sinne bezeichnen.

1.3. Stellung der Informationstheorie im statistischen Modell

Der Bereich der Informationstheorie, die sich mit dem Wesen der Entstehung und Übertragung von Informationen befasst, entspricht fast vollständig dem Bereich der von de Saussure geforderten Semiologie¹.

¹ De Saussure, a. a. O., S. 19.

Den grundlegenden Begriff dieser Theorie bildet die selektive Information, deren Definition von C. Shannon stammt¹. In der Formulierung von Neidhardt „die 'selektive Information' einer Nachricht ist die Grösse, die den Freiheitsgrad angibt, der dem Sender gelassen ist um die zu übertragende Nachricht aus der Summe möglicher Nachrichten auszuwählen“². Sie wird mit dem negativen Wert des Logarithmus der Wahrscheinlichkeit der Nachricht gemessen. Da die selektive Information eine Funktion der Wahrscheinlichkeit ist, kann sie mit dem sprachlichen Wert des sprachlichen Zeichens gleichgestellt werden.

Die Menge der Information, die in einer Nachricht enthalten ist, ist reziprok zur Wahrscheinlichkeit dieser Nachricht. Selbstverständlich dürfen auch Zeichen bzw. Symbole eines Zeichensystems als Nachrichten betrachtet werden. Im Bereich der Sprachtheorie wurde — in etwas anderer Form — die Abhängigkeit der Information von der Wahrscheinlichkeit in einem von J. Kuryłowicz formulierten semiologischen Gesetz ausgedrückt: „je enger die Verwendungssphäre des Zeichens, desto reicher sein Inhalt; je weiter die Verwendungssphäre des Zeichens, desto ärmer sein Inhalt“³.

Die selektive Information ist eine messbare Grösse. Als ihre Einheit wurde der Informationsinhalt jedes Gliedes einer Alternative angenommen, d. h. die Menge der Ungewissheit in der einfachsten Wahl zwischen zwei gleichmöglichen Nachrichten (z. B. „ja“ oder „nein“). Diese Einheit wird 'bit' genannt (abgekürzt von „binary digit“), weil sie einer Stelle im binären Zahlensystem entspricht. Im Zusammenhang damit wird die Zahl 2 als Basis des Logarithmus in der Definition der Information verwendet.

Die Anzahl der Informationseinheiten (bit) in jedem von r gleichmöglichen Symbolen gleicht dem negativen Wert des dyadi-

¹ C. Shannon, W. Weaver, *The mathematical Theory of Communication*, Urbana 1949.

² P. Neidhardt, a. a. O., S. 49.

³ J. Kuryłowicz, *Linguistique et théorie du signe*, „Journal de psychologie“, 41, No. 2 (1949), S. 172. Dass es sich hier nicht um 'Bedeutung' oder Information im populären Sinne handelt, s. S. K. Schaumjan, *Der Gegenstand der Phonologie*, „Zeitschr. f. Phonetik u. allg. Sprachwissenschaft“, 10 (1957), S. 200ff.

schen Logarithmus der Wahrscheinlichkeit des Symbols $1/r$. Bezeichnet man diese Information als H' und den dyadischen Logarithmus als ld , so kann man schreiben:

$$H' = -\text{ld } 1/r$$

oder

$$H' = -\text{ld } \bar{p} \left[\frac{\text{bit}}{\text{Symbol}} \right] \quad (1.1)$$

worin $\bar{p}$ ($= 1/r$) die durchschnittliche Wahrscheinlichkeit jedes Symbols bezeichnet.

Wenn dagegen ein System mit r Symbolen gegeben ist, von denen die Symbole $n_1, n_2, \dots, n_r$ entsprechende verschiedene Wahrscheinlichkeiten $p_1, p_2, \dots, p_r$ besitzen, d. h. jedem Symbol n_i eine Wahrscheinlichkeit p_i entspricht, dann beträgt die durchschnittliche Information pro Symbol H

$$H = -\sum_{i=1}^r p_i \text{ld } p_i \left[\frac{\text{bit}}{\text{Symbol}} \right] \quad (1.2)$$

Die durchschnittliche Information pro Symbol eines Symbolsystems wird Entropie dieses Systems genannt. Die Formel (1.1) bestimmt die maximale Entropie H' , die Formel (1.2) die Entropie H , in der die ungleichen Wahrscheinlichkeiten der Systemelemente berücksichtigt werden.

Quotient der beiden Entropien heisst relative Entropie:

$$h = \frac{H}{H'} \quad (1.3)$$

Die Ergänzung der relativen Entropie zu Eins ist ein Mass für die Redundanz des Systems:

$$R = 1 - h \quad (1.4)$$

Die Redundanz (vom lat. *redundantia*) eines Codes bezeichnet Meyer-Eppler¹ als „jene Eigenschaft, die es ermöglicht, die Kenntnis der Häufigkeitsverteilung der Zeichen zum Erraten von Teilen der Nachricht auszunutzen“.

¹ W. Meyer-Eppler, *Die Informationstheorie*, „Die Naturwissenschaften“, 49 (1952), S. 341f.

Die Übertragung der Information erfolgt in einem Kommunikationssystem, das man folgendermassen schematisch beschreiben kann¹: die Nachricht entsteht in einer Informationsquelle und wird zum Sender geschickt; der Sender codisiert sie so, dass sie in Form von Signalen über den Kanal zum Empfänger übertragen werden kann; im Empfänger werden die Signale entcodiert und die Nachricht gelangt zur Bestimmung.

Der Übertragungskanal ist ein physikalisches Medium und besitzt eine bestimmte Übertragungskapazität. C. Shannon² hat die Kanalkapazität der Übertragung für Symbole von verschiedener Dauer und Einschränkungen der Symbolfolgen folgendermassen definiert:

$$C = \lim_{T \rightarrow \infty} \frac{\log N(T)}{T} \left[\frac{\text{bit}}{\text{sek}} \right] \quad (1.5)$$

wobei $N(T)$ die Zahl der zulässigen Signale mit der Dauer T ist.

Aus dieser Definition geht hervor, dass je grösser die Zahl der Informationseinheiten (bit) ist, die in einer Zeiteinheit (sek) durch den Kanal übertragen werden können, desto grösser ist auch die Kapazität des Kanals.

Im Falle der sprachlichen Kommunikation lässt sich das Problem der Kanalkapazität auf das Problem der Wirksamkeit des angewandten Codes beschränken. Wenn man von den physikalischen Eigenschaften des Mediums, der Lautwellen usw. absieht, so hängt die Übertragungsgeschwindigkeit der Information von der Grösse der „Informationsladung“ der Symbole ab, die in einer Zeiteinheit übertragen werden können. Nimmt man an, dass den elementaren Symbolen des Codes (z. B. den Phonemen) Signale von ungefähr gleicher Dauer entsprechen (z. B. die gesprochenen Laute), dann sollte in einem wirksamen Code die Zahl der elementaren Symbole in den zusammengesetzten Symbolen, d. h. die Länge ($=$ Dauer) der Symbole (z. B. der Worte) reziprok zu der in ihnen enthaltenen Informationsmenge sein. Anders gesagt: je weniger Information, desto kürzer müsste ein Symbol sein.

Es wurde bereits oben darauf hingewiesen, dass der Informationsinhalt eines Symbols reziprok zu seiner Wahrscheinlichkeit ist;

¹ P. Neidhardt, a. a. O., S. 46.

² C. Shannon, W. Weaver, a. a. O., S. 7.

die Wahrscheinlichkeit ist ein Ausdruck für seine relative Häufigkeit — je häufiger ein Symbol vorkommt, desto weniger Information enthält es. In Verbindung mit der Formel (1.5) ergibt sich daraus folgende Forderung an einen wirksamen Code: je öfter ein Symbol vorkommt, desto kürzer sollte es sein.

Diese Forderung, den Kostenaufwand der Informationsübertragung auf ein Minimum zu senken, hat ihre theoretische Grundlage im fundamentalen Lehrsatz für den rauschfreien Kanal von Shannon gefunden¹. Dieser Lehrsatz lässt sich auch als Kriterium zur Beurteilung der Wirksamkeit eines sprachlichen Codes anwenden. Am anschaulichsten kommt es zum Ausdruck in der Darstellung, die von Herdan stammt²: als wirksam kann man einen solchen Code bezeichnen, in welchem die zeitliche Dauer (d. h. die Länge) $x_1, x_2, \dots, x_r$ der Symbole $n_1, n_2, \dots, n_r$ reziprok zu ihren Häufigkeiten $p_1, p_2, \dots, p_r$ ist. Im allgemeinen gilt:

$$x_i = f(1/p_i)$$

Für die Funktion $f(1/p_i)$ kann man eine logarithmische Funktion mit der Basis 2 wählen; dann gilt:

$$x_i = \lg(1/p_i) = -\lg p_i$$

In einem wirksamen Code muss die Grösse $(-\lg p_i)$ eine gute Annäherung an x_i sein; man kann daher schreiben:

$$-\lg p_i \leq x_i < 1 + (-\lg p_i)$$

Setzt man für $(-\lg p_i)$ den durchschnittlichen Wert, so erhält man die Entropie H . Der durchschnittliche Wert von x_i entspricht der mittleren Länge der Codesymbole M (s. die Formel 1.9. unten). Die Ungleichung lässt sich dann folgendermassen schreiben:

$$H \leq M < 1 + H \quad (1.6)$$

Aus dieser Ungleichung kann ein Index der Wirksamkeit des Codes gebildet werden:

$$c = H/M \leq 1 \quad \left[\frac{\text{bit}}{\text{Symbol}} : \frac{\text{sek}}{\text{Symbol}} = \frac{\text{bit}}{\text{sek}} \right] \quad (1.7)$$

¹ C. Shannon, W. Weaver, a. a. O., S. 29.

² G. Herdan, a. a. O., S. 171.

Die Grösse c bezeichnet, im Grunde genommen, die Kanalkapazität, wobei der Code selbst als ein Kanal betrachtet wird. $H = M$ gilt nur in einem idealen Code, in dem alle Möglichkeiten der Informationsübertragung ausgenutzt werden; in diesem Falle wird $c = 1$ und die Übertragungsgeschwindigkeit erreicht ihr Maximum C/H . In allen anderen Fällen wird c kleiner als Eins sein — das bedeutet, dass der Code in diesem Falle „schlechter“ ist und dass die Übertragungsgeschwindigkeit abnimmt. Die Wirksamkeit des Codes kann also durch die Grösse c bestimmt werden, deren zahlenmässiger Betrag zwischen Null (absolute Unwirksamkeit) und Eins (absolute Wirksamkeit) liegt. Nicht jeder Code jedoch, dessen Wirksamkeitsindex grösser als Null ist, kann als wirksam bezeichnet werden.

Die Grenzen, zwischen denen sich das Mass der Wirksamkeit bewegen kann, werden durch die Ungleichung (1.6) enger gefasst. Diese Ungleichung lässt sich folgendermassen schreiben:

$$1 \Rightarrow H/M > \frac{H}{1+H} > 0 \quad (1.8)$$

Das bedeutet, dass in einem wirksamen Code, in dem die Menge der übertragenen Information im günstigen Verhältnis zum Kostenaufwand steht, der Wirksamkeitsindex nicht nur grösser als 0, sondern auch grösser als $\frac{H}{1+H}$ sein muss.

Die Grösse

$$\frac{H}{1+H} = c'$$

wird in der vorliegenden Untersuchung als obere Unwirksamkeitsgrenze bezeichnet und zusammen mit der Grösse c zur Charakterisierung der sprachlichen Code verwendet.

1.4. Plan der Untersuchung

Die linguistische Grundlage der Untersuchung bildet grundsätzlich die Sprachauffassung von de Saussure. In folgenden Punkten weicht die dieser Untersuchung zugrundeliegende Sprachauffassung von der von de Saussure ab: 1) die Sprache und das

Sprechen werden als zwei untrennbare und voneinander abhängige Seiten der menschlichen Rede betrachtet; 2) es wird die Möglichkeit und Zweckmässigkeit einer panchronischen Untersuchung der Sprache vorausgesetzt¹.

Diese beiden Punkte hängen miteinander zusammen: die statistische Auffassung der menschlichen Rede lässt aus dem Sprechen Daten gewinnen, die die Sprache charakterisieren; eine Zusammenstellung der Daten, die sich auf mehrere aufeinanderfolgende synchronische Sprachzustände ein und derselben Sprache beziehen, schafft ein panchronisches Bild dieser Sprache.

Die vorliegende Untersuchung befasst sich mit drei Sprachzuständen der persischen Sprache. Die Untersuchung jedes Sprachzustandes wird folgendermassen durchgeführt:

a) dem Kollektiv des Sprachzustandes werden Proben (Texte) entnommen;

b) unter Berücksichtigung der statistischen Sicherheit werden auf Grund dieser Proben die das Kollektiv charakterisierenden Verteilungen bestimmt;

c) jeder Sprachzustand wird als ein Code betrachtet, der aus zusammengesetzten und elementaren Symbolen besteht;

d) die Wörter werden als zusammengesetzte Symbole betrachtet, die aus zweierlei elementaren Symbolen gebildet sind: 1. aus Silben, 2. aus Phonemen;

e) die elementaren Symbole jeder von den beiden Arten werden als gleichlange Signale betrachtet; damit wird die Dauer (die Länge) der Wörter zweifach bestimmt: 1) nach der Silbenzahl, 2) nach der Phonemzahl.

Die Wörter werden nach den zwei Unterscheidungsmerkmalen folgendermassen aufgeteilt:

a) in Worttypen — nach der Anzahl der Silben; z. B. Einsilber, Zweisilber, usw.;

b) in Wortklassen — nach der Anzahl der Laute (die den Phonemen entsprechen); z. B. Einlauter, Zweilauter; usw.

¹ Diese Voraussetzung steht in Übereinstimmung mit Forderungen mancher Vertreter des Strukturalismus, s. z. B. V. Brøndal, *Linguistique structurale*, „Acta Linguistica“, 1 (1939), S. 4ff.

Die Worte mit einer bestimmten Anzahl von Silben und Lauten werden Wortgruppen genannt, z. B. Einsilber-Einlauter, Einsilber-Zweilauter, usw.)¹.

Jeder Sprachzustand wird als eine Verteilung von Worttypen und eine von Wortklassen untersucht. Die gegenseitige Beziehung zwischen Silben- und Lautzahl in den Wortgruppen wird als statistische Korrelation charakterisiert.

Es werden folgende Charakteristika ermittelt:

1) Entropie, maximale Entropie, relative Entropie und Redundanz der Wörter nach der Silbenzahl — H_s , H'_s , h_s , R_s ;

2) Entropie, maximale Entropie, relative Entropie und Redundanz der Wörter nach der Phonemzahl — H_{ph} , H'_{ph} , h_{ph} , R_{ph}

3) Wirksamkeitsindices und die oberen Unwirksamkeitsgrenzen in bezug auf die Silbenzahl — c_s , c'_s und in bezug auf die Phonemzahl — c_{ph} , c'_{ph} .

Zu den informationstheoretischen kommen noch rein statistische Charakteristika hinzu:

1) mittlere Wortlänge und die mittlere Abweichung

a) nach der Silbenzahl — M_s , σ_s ,

b) nach der Phonemzahl — M_{ph} , σ_{ph} ,
nach den Formeln²:

$$M = \sum_{i=1}^r p_i x_i \quad (1.9)$$

$$\sigma = \sqrt{\sum_{i=1}^r p_i (x_i - M)^2} \quad (1.10)$$

Darin bedeutet: p_i — Wahrscheinlichkeit des Worttyps bzw. der Wortklasse mit 1...r Silben bzw. Phonemen;
 x_i — die Anzahl der Unterscheidungsmerkmale, also 1...r Silben bzw. Phoneme.

¹ Diese Klassifizierung stammt von P. Menzerath, W. Meyer-Eppler, *Sprachtypologische Untersuchungen*, „Studia Linguistica“, 4 (1950), S. 54—93.

² Vgl. G. Herdan, a. a. O., S. 306 ff.

- 2) mittlere Silbenlänge nach der Phonemzahl — M_{ph}/M_s ;
 3) Koeffizient der Korrelation zwischen der Silbenzahl und der Phonemzahl im Wort — r , nach der Formel¹

$$r = \frac{\sum xy - n \bar{x} \bar{y}}{n \sigma_x \sigma_y} \quad (1.11)$$

Darin bedeutet:

- $\bar{x}$ — die mittlere Wortlänge in Silben,
 $\bar{y}$ — die mittlere Wortlänge in Phonemen,
 n — die Anzahl der untersuchten Fälle,
 $\sum xy$ — die Summe der Produkte der koordinierten Werte von x und y , mal die Zahl der durch x , y , charakterisierten Wortgruppen; z. B. bei 50 Wortgruppen: Einsilber-Dreilauter $\sum xy = 50 \cdot (1 \cdot 3) = 150$.

Der Betrag von r bezeichnet den Grad der Korrelation, der sich zwischen zwei Grenzwerten, dem funktionellen Zusammenhang bei $r = |1|$, und der Unabhängigkeit der Variablen bei $r = 0$, bewegen kann.

2. SPRACHLICHES MATERIAL

2.1. Aufteilung des sprachlichen Materials: Wort, Silbe, Phonem

2.1.1. Das Wort

In der vorliegenden Untersuchung, die sich ausschliesslich mit dem Ausdrucksplan der Sprache befasst, wird unter dem Begriff „Wort“ ein Symbol, das von anderen Symbolen durch eine Pause getrennt werden kann, verstanden.

Vom linguistischen Standpunkt aus ist diese Formulierung des Wortbegriffes nicht ausreichend, weil sie sich nur auf eine Seite des sprachlichen Zeichens — auf das „signifiant“ — bezieht; darüber hinaus bezieht sie sich auf ein Produkt einer schon vorhandenen Aufteilung des Redeflusses, die sich ihrerseits unter

¹ G. Herdan, a. a. O., S. 350f.

Berücksichtigung von beiden Seiten des sprachlichen Zeichens vollzog.

Aus diesem Grunde muss der Begriff des Wortes auch die signifié-Seite umfassen, um als Kriterium bei der Aufteilung des sprachlichen Materials, auf dem die vorliegende Untersuchung durchgeführt wird, dienen zu können.

Die Formulierung dieses Begriffes muss genügend allgemein gehalten sein, damit sie sich auf alle untersuchten Sprachen anwenden lässt. Wie A. S. C. Ross festgestellt hat, ist es unmöglich, eine präzise und gleichzeitig für alle Sprachen geltende Definition des Wortes anzugeben¹.

In der vorliegenden Arbeit bezeichnet der Begriff „Wort“ eine Verbindung eines kleinsten selbständigen Elementes des Ausdrucksplanes — eines selbständigen Kenem, mit einem kleinsten selbständigen Element des Inhaltsplanes — einem selbständigen Plerem.

Diese Formulierung entspricht im grossen und ganzen der Definition der sprachlichen Einheit bei de Saussure:

„Die Einheit (ist) ... eine Lautfolge, welche mit Ausschluss des in der gesprochenen Reihe Vorausgehenden und Darauffolgenden das Bezeichnende für eine gewisse Vorstellung ist“². Die Bezeichnung „kleinstes selbständiges Element“ in der Formulierung entspricht der Bezeichnung „kleinste freie Form“ in der Terminologie von L. Bloomfield³.

Die in der vorliegenden Untersuchung angenommene Formulierung des Wortbegriffes dient nur als ein allgemeines Kriterium bei der Aufteilung der untersuchten Texte. Daneben gilt das Prinzip, dass gleiche grammatische Kategorien in allen drei untersuchten Sprachen gleich behandelt werden.

In allen drei untersuchten Sprachen lassen sich solche Wörter bzw. Worтеlemente feststellen, die allein stehen können, und solche, die nur in Verbindung mit anderen auftreten.

¹ Alan S. C. Ross, *The fundamental definitions of the theory of language*, „Acta Linguistica“, Bd. IV (1944) S. 105.

² F. de Saussure, a. a. O., S. 124.

³ Leonard Bloomfield, *A Set of Postulates for the Science of Language*, „Journal of the Ling. Society of America“, Bd. II (1926) S. 175.

Die ersteren sind rein lexikalische Formen. Die Frage nach ihrer Selbständigkeit wird in der vorliegenden Untersuchung für jede Sprache einzeln an Hand ihres Vorkommens in den unten verzeichneten Wörterbüchern entschieden.

Die letzteren — die man als grammatische Affixe bezeichnen kann¹ — können je nach dem angewandten Kriterium, sowohl als selbständige, als auch nicht selbständige Wörter betrachtet werden.

Die Aufteilung des Altpersischen in Wörter in der vorliegenden Untersuchung folgt der altpersischen Schreibung²; damit wird auch die Frage der altpersischen Affixe entschieden. Über die Behandlung der mittelpersischen und neupersischen Affixe s. unten.

Der Wortakzent wird in der vorliegenden Untersuchung nicht berücksichtigt, weil die Stellung und Rolle des Akzents im Alt- und Mittelpersischen unbekannt ist.

Im Neupersischen hat der Akzent nur eine morphologische und keine phonologische Bedeutung, d. h. er fungiert als Unterscheidungsmerkmal nur in den Wörtern, die ein und derselben morphologischen Familie angehören³. Aus diesem Grunde kann auch der neupersische Wortakzent in der vorliegenden Untersuchung ausser acht gelassen werden.

2.1.2. Die Silbe

Die vorliegende Untersuchung setzt sich als Ziel, u. a. die Rolle der Silbe in der sprachlichen Kommunikation zu untersuchen; dabei werden die phonetischen und phonologischen Probleme des Silbenbaues ausser acht gelassen.

Unter dem Begriff „Silbe“ wird demnach ein Gebilde verstanden, das aus einem oder mehreren elementaren Symbolen (Phonemen) besteht und in dem ein elementares Symbol der Silbenkern ist.

¹ *Handbuch der Orientalistik*, hgb. von B. Spuler, Erste Abt., IV Bd. *Iranistik*, I Abschnitt — *Linguistik*; Leiden—Köln, 1958, S. 190.

² Roland G. Kent, *Old Persian*, New Haven, Conn. 1953 S. 19, § 44.

³ Jiří Krámský, *A Study in the Phonology of Modern Persian*, „Archiv Orientální“, Bd. 11, 1939, S. 80.

Als Silbenkerne werden kraft Voraussetzung die Vokale und Diphthonge betrachtet.

Demnach wird die Wortlänge in Silben nach der Anzahl der Silbenkerne im Wort bestimmt.

2.1.3. Das Phonem

Unter dem Begriff „Phonem“ wird in der vorliegenden Untersuchung das elementare Symbol verstanden, aus dem die zusammengesetzten Symbole gebildet werden.

Vom linguistischen Standpunkt aus liegt dieser Formulierung die Definition dieses Begriffes von Trubetzkoy zugrunde:

„Phonologische Einheiten, die sich vom Standpunkt der betreffenden Sprache nicht in noch kürzere aufeinanderfolgende phonologische Einheiten zerlegen lassen, nennen wir Phoneme. ... Man darf sagen, dass das Phonem die Gesamtheit der phonologisch relevanten Eigenschaften eines Lautgebildes ist“¹.

2.2. Altpersisch

Unter dem Begriff „Altpersisch“ wird in der vorliegenden Untersuchung die Sprache der Keilinschriften der achämenidischen Könige — ein südwestiranischer Dialekt — verstanden.

Die hier untersuchten Texte stammen aus der Zeit der Herrschaft von Dareios dem Grossen (521 — 486) und Xerxes (486 — 465); sie stellen den Stand der persischen Sprache im VI—V Jhh. v. d. Z. dar.

2.2.1. Phonologisches System

In der vorliegenden Untersuchung wird das phonologische System des Altpersischen nach R. G. Kent² angenommen:

Vokale und Diphthonge:	<i>a i u ā ī ū r</i>
	<i>ai au āi āu</i>
Halbvokale:	<i>y v</i>
Liquide:	<i>r l</i>

¹ N. S. Trubetzkoy, *Grundzüge der Phonologie*, „Travaux du Cercle linguistique de Prague“, VII, 1939 S. 34ff.

² a. a. O., S. 25, § 59.

Nasale:	<i>m n</i>
Gutturale:	<i>k g χ</i>
Palatale:	<i>č ĵ</i>
Dentale:	<i>t d ð</i>
Labiale:	<i>p b f</i>
Sibilante:	<i>s š ʃ z</i>
Aspirate:	<i>h</i>

2.2.2. Umschrift in Phoneme

Die sog. normalisierte Umschrift der altpersischen Inschriften, wie sie von Kent durchgeführt wird, ist keine rein phonemische Umschrift¹.

Die wichtigsten Unterschiede zwischen der normalisierten und der phonemischen Umschrift, die bei der Vorbereitung der Texte zur Untersuchung berücksichtigt wurden, sind die folgenden:

a) normalisiertes finales *-ā*, *-īy*, *-ūv* repräsentiert phonemisches *-a*, *-ī*, *-ū*²; z. B.

normalisiert	phonemisch
<i>utā</i> „und“	<i>uta</i>
<i>ðātiy</i> „es spricht“	<i>ðāti</i>
<i>hauv</i> „er, jener“	<i>hau</i>

b) normalisiertes *-īy*-, *-uv*- repräsentiert postkonsonantisches *-y*-, *-v*-; z. B.

normalisiert:	phonemisch:
<i>xšāyadīya</i> „der König“	<i>xšāyadya^h</i>
<i>haruva-</i> „sämtlicher“	<i>harva-</i>

c) normalisiertes anlautendes *u*- repräsentiert phonemisches *^hu*-; normalisiertes *uv*- repräsentiert phonemisches *^hv*- oder *^huv*-; z. B.

normalisiert:	phonemisch:
<i>utava</i> „kräftig“	<i>^hutava^h</i>

¹ a. a. O., S. 19f, § 45.

² a. a. O., S. 19, § 45 Fn..

d) eine Reihe von Wörtern verbirgt hinter der normalisierten Schreibung *-ar*- den ap. *r*-Vokal¹; z. B.

normalisiert:	phonemisch:
<i>karta-</i> „getan“	<i>krta-</i> , vgl. skr. <i>krtá</i> , av. <i>karəta-</i>

e) jedes normalisierte finale *-a* repräsentiert ein *-a* plus einen ungeschriebenen finalen Konsonanten²; z. B.

normalisiert:	phonemisch:
<i>abara</i>	<i>abara^t</i> „er brachte“ vgl. skr. <i>ábharat</i>
<i>hya</i> „welcher“	<i>hya^h</i> , vgl. skr. <i>syás</i>
<i>tya</i> „welches“	<i>tya^d</i> , vgl. skr. <i>tyád</i>

f) manchmal ein normalisiertes finales *-ā* repräsentiert ein *-ā* plus einen ungeschriebenen finalen Konsonanten³, z. B.

normalisiert:	phonemisch:
<i>napā</i> „Enkel“	<i>napā^t</i> , vgl. skr. <i>napāt</i>
<i>Pārsā</i> „in Pārsā“	<i>Pārsā^d</i> , vgl. skr. Ablativ-Endung <i>-ād</i> ;

g) die vorkonsonantischen Nasale wurden in der ap. Schrift weggelassen und so auch in der normalisierten Umschrift⁴; z. B.

normalisiert:	phonemisch:
<i>Kabūjiya</i> „Kambyses“	<i>Kambūjya^h</i>
<i>hātīy</i> „sie sind“	<i>hanti</i> vgl. skr. <i>sānti</i> .

2.3 Mittelpersisch

Unter dem Begriff „Mittelpersisch“ wird in der vorliegenden Untersuchung ausschliesslich das Buchpahlavi verstanden, d. h. die mittelpersische Schriftsprache, in der die zarathustrischen Schriften religiösen und profanen Inhalts in der Zeit zwischen dem 6 und 10 Jh. verfasst worden sind.

¹ a. a. O., S. 15f, § 30.

² a. a. O., S. 17, § 36; S. 18, § 40.

³ a. a. O., §§ 36, 40 und Lexicon, S. 164ff.

⁴ a. a. O., S. 17f, § 39.

Die hier untersuchten Texte sind den profanen Werken unterhaltenden Charakters entnommen (*Aβiyātkār ī Zarērān*, *Kārnamak ī Artaxšēr*, *Mātikān ī Čatrang*). Ihre Entstehungszeit lässt sich nicht mit Sicherheit bestimmen¹; auf jeden Fall stellt ihre Sprache die dem Neupersischen vorangehende Entwicklungsstufe dar.

2. 3. 1. Phonologisches System

Das Problem des phonologischen Systems des Mittelpersischen ist unmittelbar mit dem Problem der Lesung der mittelpersischen Texte verbunden. Es existiert keine allgemein anerkannte Umschrift des Mittelpersischen².

In der vorliegenden Untersuchung wird das Transkriptionssystem von H. F. J. Junker angenommen, wie es von ihm zur Wiedergabe der Heterogrammlösungen in dem *Frahang ī Pahlavik*³ angewendet wird.

Der diesem Transkriptionssystem zugrundeliegende Phonembestand stimmt grundsätzlich mit dem von Salemann im „Grundriss der iranischen Philologie“ S. 255—75 angenommenen überein.

Dieses System umfasst folgende Phoneme:

Vokale:	<i>a i u ā ē ī ō ū</i>
Halbvokale:	<i>y v</i>
Liquide:	<i>r l</i>
Nasale:	<i>n m</i>
Gutturale:	<i>k g x xʷ γ</i>
Palatale:	<i>č ĵ</i>
Dentale:	<i>t d θ δ</i>
Labiale:	<i>p b f β</i>
Sibilante:	<i>s z š ž</i>
Aspirate:	<i>h</i>

¹ J. C. Tavadia, *Die mittelpersische Sprache und Literatur der Zarathustrier*, Iranische Texte u. Hilfsbücher hgb. von H. F. J. Junker, Leipzig 1956, S. 135—140.

² *Handbuch d. Orientalistik* hgb. von B. Spuler, *Iranistik I* Abschnitt „Linguistik“, Leiden—Köln 1958, S. 121.

³ Heinrich F. J. Junker, *Das Frahang ī Pahlavik*, Leipzig 1955.

2.3.2. Die Aufteilung in Wörter

Der Aufteilung der mittelpersischen Texte in Wörter liegen in der vorliegenden Untersuchung folgende Wörterbücher zugrunde:

Nyberg H. S. — „Hilfsbuch des Pehlevi“, Bd. II, Uppsala 1931

Kapadia D. D. — „Glossary of Pahlavi Vendidad“, Bombay 1953

Folgende grammatische Affixe werden als selbständige Wörter betrachtet:

- a) beim Nomen: *ī-īdāfat*,
die Dativpostposition *rā*
die Präpositionen: *pat* „zu“, „mit“, „für“,
ō „zu“, „für“;
- b) beim Verbum: die Partikel des vollendeten Aspektes *bē*
die Verbalpräposition *hamē*
der Negationspräfix *nē*

2.4. Neupersisch

Unter dem Begriff „Neupersisch“ wird in der vorliegenden Untersuchung die heutige Litteratursprache Irans verstanden, deren älteste Sprachdenkmäler aus dem 10. Jh. d. Z. stammen.

Die hier untersuchten Texte sind den Werken moderner iranischer Schriftsteller entnommen und geben den Sprachzustand des 20. Jh. wieder.

2.4.1. Phonologisches System

In der vorliegenden Untersuchung wird ein phonologisches System des Neupersischen nach J. Krámský¹ und M. Shaki² angenommen.

¹ Jiří Krámský, *A Study in the Phonology of Modern Persian*, „Archiv Orientální“, Bd. 11 (1939), S. 66—83

² Mansour Shaki, *The Problem of the Vowel Phonemes in the Persian language*, „Archiv Orientální“, Bd. 25 (1957).

Dieses System umfasst folgende Phoneme:

Vokale:	<i>a e o ā i u</i>
Halbvokale:	<i>y (w)</i>
Liquide:	<i>r l</i>
Nasale:	<i>n m</i>
Gutturale:	<i>k g x γ</i>
Palatale:	<i>č ĵ</i>
Dentale:	<i>t d</i>
Labiale:	<i>p b f</i>
Labiodentale:	<i>v</i>
Sibilante:	<i>s z š ž</i>
Aspirate:	<i>h</i>

In dem hier angenommenen System kommen keine Diphtonge vor. Die Frage der neupersischen Diphtonge ist strittig¹; nach der Ansicht von Krámský² sind die Lautgebilde: *ai*, *au*, *āi*, *āu*, in Wörtern wie: *nai* „Flöte“, *nau* „neu“, *ĵai* „Platz“, *muĵi* „Haar“, zerlegbare Lautverbindungen vom Typ: Vokal + Halbvokal (z. B. *nai* „eine Flöte“ — *na-i* „irgendeine Flöte“, *rau* „geh“ — */mi/ra-vim* „wir gehen“; usw.) und lassen sich deshalb mit den 6 Regeln für die monophonematische Wertung von Trubetzkoy³ nicht in Einklang bringen.

Das konsonantische *y* geschr. *w*, wird hier als eine Stellungsvariante des labiodentalen Konsonanten *v* betrachtet; vgl. sein Verhalten in den arabischen Lehnwörtern, z. B.:

sing. <i>waqt</i> „Zeit“	— [<i>vaqt</i>]
plur. <i>awqāt</i>	[<i>auqāt</i>]

Die arabischen Lehnwörter sind dem neupersischen phonologischen System angepasst⁴, d. h.:

a) die Opposition: emphatisch — nicht-emphatisch, ist aufgehoben;

¹ *Handbuch der Orientalistik* a. a. O., S. 183f.

² J. Krámský, a. a. O., 71f.

³ N. S. Trubetzkoy, *Anleitung zu phonologischen Beschreibungen*, Brno 1935, S. 11f.

⁴ J. Krámský, a. a. O., S. 72ff.

b) die glottalen Plosiven 'Ain und Hamza haben ihre Unterscheidungsfunction verloren und werden nur in der Schrift beibehalten.

2.4.2. Umschrift in Phoneme

Die Umschreibung der neupersischen Texte in Phoneme erfolgt in der vorliegenden Arbeit nach den Lesungen von V. Mil'ler in seinem *Persidsko—russkij slovar*, Moskwa 1953.

Dabei wird *ī-iḏāfat* als kurzes *e* umgeschrieben, ungeachtet des Auslautes des vorangehenden Wortes, z. B.:

<i>ketāb e man</i>	„mein Buch“
<i>mu e man</i>	„mein Haar“

Das Bindewort „und“ ist in der arabischen Schrift ein Homograph. Es kann sowohl das persische: *o* (<ir. *uta*), als auch das arabische LW: *va* repräsentieren.

In der vorliegenden Untersuchung wird es in folgenden Stellungen als *o* gelesen:

- in idiomatischen Ausdrücken vom Typ:
āmad o raft — „Das Kommen und Gehen; Betrieb“,
- zwischen zwei Synonymen; z. B.:
tašwiš-o-eḏterāb — „Aufregung“
- zwischen zwei verbalen Sätzen vom Typ:
mi āmad o mi goft — „er kam und sagte“
- zwischen zwei Nomina, die eng miteinander verbundene Objekte bezeichnen z. B.:
miz o šandali — „Tisch und Stuhl“

In allen anderen Fällen wird „und“ als *va* gelesen.

2.4.3. Aufteilung in Wörter

Die folgenden grammatischen Affixe werden als selbständige Wörter betrachtet:

- beim Nomen: *ī-iḏāfat*

der Dativpräfix *be*; (ausgenommen sind die festen Zusammensetzungen: *bedin*, *bedān*); die Akkusativpostposition *rā* (ausgenommen ist die Zusammensetzung: *čerā*);

- b) beim Verbum: die Konjunktivpartikel *be*
 die Verbalpräposition *mi*
 der Negationspräfix *na, ma*.

2.5. Gewinnung des Materials

Das statistische Material, mit dem die Untersuchung durchgeführt wird, ist den Textproben entnommen.

Das Mittel- und Neupersische, d. h. die Sprachen, die genügend Material bieten, liefern je drei Proben, von denen jede von einem anderen Verfasser stammt. Auf diese Weise wurde versucht, die stilistische Mannigfaltigkeit der literarischen Sprache zu berücksichtigen.

Mit der Ausnahme von *Mātikān ī Čatrang*, dessen Text nur etwa 900 Wörter enthält, besteht jede von diesen Proben wieder aus drei kleineren Teilproben. Auf diese Weise konnte mit Hilfe des χ^2 — Testes¹ die innere Einheitlichkeit der Proben, die ein und demselben Verfasser entnommen wurden, gesichert werden.

Diesem Schema wurde die Einteilung des altpersischen Materials angepasst. Es sind nur wenige altpersische Texte vorhanden, die für eine statistische Untersuchung genügend lang sind. Die folgenden wurden untersucht: die Behistun-Inschrift von Dareios und die sogen. Daiva-Inschrift von Xerxes. Die Texte sind zu kurz um sie in kleinere Proben zerteilen zu können. Es wurde deshalb notwendig, auf die Prüfung der inneren Einheitlichkeit der altpersischen Proben zu verzichten.

Insgesamt dienten folgende Texte als Material:

A. Altpersisch:

1. D. I—V.: Die Tafeln I, II, IV, V, der Behistun-Inschrift von Dareios; s. Kent, S. 116—122, 128—133. Die Formel: *θātiy Dārayavauš xšāyadīya* wurde in jeder Tafel nur einmal einbezogen.

Umfang: 2249 Worte

¹ S. u. a. G. Herdan, a. a. O., S. 88ff.

2. D. IV.: Die Tafel IV der Behistun-Is., s. Kent, S. 128—130.

Umfang: 619 Worte

3. X.: Die Daiva-Is. von Xerxes; s. Kent, S. 150—151.

Umfang: 442 Worte

B. Mittelpersisch:

1. Kn.: *Kārnāmak ī Artaxšēr ī Pāpakān*, die ersten drei Kapitel, in: Nyberg, *Hilfsbuch des Pehlevi* Bd. I.

Umfang: 1985 Worte

2. Čn.: *Mātikān ī Čatrang*, in: *The Pahlavi Texts* von Khudayar D. Sh. Irani, Bombay 1899.

Umfang: 892 Worte

3. AZ.: *Aβiyātkār ī Zarērān*, in: *Pahlavi Texts* von Jamaspji M. Jamasp-Asana Bd. I. Bombay 1897.

Umfang: 1953 Worte

C. Neupersisch:

1. Hed.: *Šādeq-e Hedāyat, Hāji āqā*, Teheran 1334 H. š.

Umfang: 2567 Worte

2. Mas.: *Moḥammad-e Mas'ūd*, drei Werke:
Dar talāš-e ma'āš o. O. 1323 H. š.
Ašraf-e maḫluqāt o. D. Teheran
Tafriḫāt-e šab o. O. 1323 H. š.

Umfang: 2475 Worte

3. Hej.: *Moḥammad-e Hejāzi, Sāyar*, Teheran 1332 H. š.

Umfang: 2578 Worte

Die Texte wurden phonemisch umgeschrieben.

Das Zählen erfolgte mit Hilfe von Zähllisten für Wortgruppen nach der Methode von Menzerath und Mayer-Eppler¹.

¹ P. Menzerath, W. Meyer-Eppler, a. a. O., S. 60f.

Die Untersuchung der Einheitlichkeit der Proben wurde nur für das Mittel- und Neupersische durchgeführt.

Es stellte sich heraus, dass eine aus drei Teilproben bestehende Probe von ungefähr 2000—2500 Wörtern die für einen Verfasser charakteristischen Worttypen- und Wortklassenverteilungen mit ziemlicher Genauigkeit erfassen lässt.

Jeder Sprachzustand der persischen Sprache wird in der vorliegenden Untersuchung durch drei Proben mit jeweils etwas voneinander unterschiedlichen Verteilungen repräsentiert. Demzufolge werden die einzelnen informationstheoretischen und statistischen Charakteristika jeweils als drei Werte erscheinen.

Darunter wird ein arithmetisches Mittel aus allen drei Daten gezogen. Da die drei Werte, die jedes Charakteristikum repräsentieren, die Eigenschaften statistisch unterschiedlicher Proben zum Ausdruck bringen, kann das aus ihnen gezogene arithmetische Mittel nur der Orientierung dienen; es lässt sich nicht mit Sicherheit als ein allgemeinsprachliches Charakteristikum ansehen. Statistisch gesichert sind nur die Werte, die die einzelnen Proben charakterisieren. In der vorliegenden Untersuchung weichen jedoch die drei Werte eines Charakteristikums nur wenig voneinander ab. Das arithmetische Mittel kann daher als eine gute Annäherung an ein allgemeinsprachliches Charakteristikum betrachtet werden.

3. ZAHLENMATERIAL

Das Zahlenmaterial wird für jeden der drei Sprachzustände in drei Tafeln in folgender Anordnung angeführt:

- die informationstheoretischen Charakteristika der Wortklassen,
- die informationstheoretischen Charakteristika der Worttypen,
- die statistischen Charakteristika der Wortklassen, -typen und -gruppen.

Die maximale Entropie H' für Wortklassen und Worttypen wurde auf Grund der grössten beobachteten Phonem- bzw. Silbenzahl pro Wort berechnet. Diese beträgt: a) für Wortklassen: Ap. — 14, Mp. — 16, Np. — 14; b) für Worttypen: Ap., Mp., Np. — 6.

3.1. Altpersisch

Tafel Ia

Probe	H_{ph}	H'_{ph}	h_{ph}	R_{ph}	c_{ph}	c'_{ph}
D. I—V	3,13	3,81	0,82	0,18	0,52	0,76
D. IV	3,12	3,81	0,82	0,18	0,54	0,76
X.	2,98	3,81	0,78	0,22	0,49	0,75
Mittelwert	3,08	3,81	0,81	0,19	0,52	0,76

Tafel Ib

Probe	H_s	H'_s	h_s	R_s	c_s	c'_s
D. I—V	2,02	2,59	0,78	0,22	0,76	0,67
D. IV	1,91	2,59	0,74	0,26	0,75	0,66
X.	1,99	2,59	0,77	0,23	0,75	0,67
Mittelwert	1,97	2,59	0,76	0,24	0,75	0,67

Tafel Ic

Probe	M_{ph}	σ_{ph}	M_s	σ_s	M_{ph}/M_s	r
D. I—V	5,99	2,31	2,65	1,08	2,26	0,97
D. IV	5,78	2,35	2,56	1,16	2,26	0,97
X.	6,13	2,41	2,67	1,10	2,29	0,89
Mittelwert	5,97	2,36	2,61	1,11	2,27	0,94

3.2. Mittelpersisch

Tafel IIa

Probe	H_{ph}	H'_{ph}	h_{ph}	R_{ph}	c_{ph}	c'_{ph}
Kn.	3,07	4,00	0,77	0,23	0,75	0,75
Čn.	3,17	4,00	0,79	0,21	0,71	0,76
AZ.	2,99	4,00	0,75	0,25	0,74	0,75
Mittelwert	3,08	4,00	0,77	0,23	0,73	0,75

Tafel IIb

Probe	H _s	H' _s	h _s	R _s	c _s	c' _s
Kn.	1,54	2,59	0,59	0,41	0,94	0,61
Čn.	1,65	2,59	0,64	0,36	0,94	0,62
AZ.	1,41	2,59	0,54	0,46	0,90	0,59
Mittelwert	1,53	2,59	0,59	0,41	0,93	0,61

Tafel IIc

Probe	M _{ph}	σ _{ph}	M _s	σ _s	M _{ph} /M _s	r
Kn.	4,19	2,30	1,63	0,85	2,52	0,90
Čn.	4,44	2,61	1,75	0,98	2,54	0,81
AZ.	4,03	2,17	1,56	0,72	2,58	0,91
Mittelwert	4,22	2,36	1,65	0,85	2,55	0,87

3.3. Neupersisch

Tafel IIIa

Probe	H _{ph}	H' _{ph}	h _{ph}	R _{ph}	c _{ph}	c' _{ph}
Hed.	2,96	3,81	0,78	0,22	0,81	0,75
Mas.	2,97	3,81	0,78	0,22	0,79	0,75
Hej.	2,85	3,81	0,75	0,25	0,78	0,74
Mittelwert	2,93	3,81	0,77	0,23	0,79	0,75

Tafel IIIb

Probe	H _s	H' _s	h _s	R _s	c _s	c' _s
Hed.	1,54	2,59	0,59	0,41	0,95	0,61
Mas.	1,57	2,59	0,61	0,39	0,95	0,61
Hej.	1,47	2,59	0,57	0,43	0,92	0,59
Mittelwert	1,53	2,59	0,59	0,41	0,94	0,60

Tafel IIIc

Probe	M _{ph}	σ _{ph}	M _s	σ _s	M _{ph} /M _s	r
Hed.	3,67	2,22	1,62	0,86	2,26	0,90
Mas.	3,78	2,28	1,67	0,86	2,27	0,88
Hej.	3,64	2,07	1,59	0,91	2,29	0,88
Mittelwert	3,70	2,19	1,63	0,88	2,27	0,89

4. AUSWERTUNG DER ERGEBNISSE

Eine Zusammenstellung der mittleren Werte von c_{ph} , c'_{ph} , c_s , c'_s für alle drei Entwicklungsstufen (s. Tafel IV) lässt einige interessante Züge der persischen Sprachentwicklung erkennen.

Tafel IV

	Ap.	Mp.	Np.
$\bar{c}_{ph}$	0,52	0,73	0,79
$\bar{c}'_{ph}$	0,76	0,75	0,75
$\bar{c}_s$	0,75	0,93	0,94
$\bar{c}'_s$	0,67	0,61	0,60

Man beobachtet vor allem, dass die Indices der Wirksamkeit — sowohl vom Standpunkt der Phonemzahl als auch der Silbenzahl aus — immer grösser werden. Das bedeutet, dass sich das Verhältnis der Information, die in den durch ihre Länge differenzierten Wörtern enthalten ist, zu den Übertragungskosten (der mittleren Länge nach der Phonem- und Silbenzahl) immer günstiger im Sinne des fundamentalen Lehrsatzes gestaltet.

Zwischen der Wirksamkeit vom Standpunkt der Wortlänge nach der Phonemzahl und der nach der Silbenzahl besteht jedoch ein bedeutender Unterschied: die Indices der ersteren liegen im Ap. und Mp. unterhalb der oberen Unwirksamkeitsgrenze c' , d. h. der sprachliche Code genügt im Ap. und Mp. in bezug auf die Wortlänge nach der Phonemzahl dem Kriterium der Wirksam-

keit nicht. Erst im Np. steigt der Wirksamkeitsindex ein wenig über diese Grenze. Die Indices der Wirksamkeit nach der Silbenzahl befinden sich dagegen stets oberhalb der oberen Unwirksamkeitsgrenze. Die Tatsache, dass die persische Sprache im gesamten Verlaufe ihrer Entwicklung nur in bezug auf Silben den Forderungen des fundamentalen Lehrsatzes genügt, weist darauf hin, dass in ihr die Silben, nicht aber die Phoneme den Charakter der gleichlangen Signale, d. h. der Einheiten der Übertragungskosten besitzen.

Zwei weitere Umstände scheinen die Rolle der Silbe als einer isochronen Übertragungseinheit hervorzuheben: 1) die Unveränderlichkeit ihrer Struktur im Laufe der persischen Sprachentwicklung, und 2) die starke Vergrößerung der Redundanz der Wörter von unterschiedlicher Silbenzahl.

Eine Zusammenstellung der mittleren Silbenlänge und der Koeffizienten der Korrelation der Phonem- und Silbenzahl in allen drei Entwicklungsstufen des Persischen (s. Tafel V) zeigt, dass sich die Silbe als geschlossenes Lautgebilde im Laufe der Entwicklung dieser Sprache nicht verändert hat.

Tafel V

	Ap.	Mp.	Np.
$\bar{M}_{ph}/\bar{M}_s$	2,27	2,55	2,27
$\bar{r}$	0,94	0,87	0,89

Wenn man die Redundanzwerte R_{ph} und R_s für Ap., Mp., und Np. miteinander vergleicht (s. Tafel VI) stellt man fest, dass die Vergrößerung der ersteren im Gegensatz zu der letzteren sehr gering ist.

Tafel VI

	Ap.	Mp.	Np.
$\bar{R}_{ph}$	0,19	0,23	0,23
$\bar{R}_s$	0,24	0,41	0,41

Die starke Vergrößerung der Redundanz R_s , d. h. der Möglichkeit, die Kenntnis der relativen Häufigkeit der Wörter von unterschiedlicher Silbenzahl zur Voraussage dieser Wörter auszunutzen, während dagegen die Redundanz R_{ph} fast unverändert bleibt, scheint darauf hinzuweisen, dass die erstere mit dem praktischen Gebrauch des Codes in Zusammenhang steht, während der zweiten keine praktische Bedeutung (als einem Unterscheidungsmerkmal der Wörter) zugemessen werden kann.

In diesem Zusammenhang ist eine Betrachtung der arabischen Lehnwörter im Neupersischen besonders interessant. Untersuchungen¹ zeigen, dass die arabischen Lehnwörter ungefähr 50% des neupersischen Wortschatzes und 12—20% der in den neupersischen Texten angewandten Wörter ausmachen. Aus einer Untersuchung von Fucks² ergibt sich für das Arabische: die mittlere Wortlänge nach der Silbenzahl M_s — 2,10, die Entropie H_s — 1,72. Diese beiden Werte weichen stark von den entsprechenden Werten für das Mittel- und Neupersische ab. Diese Tatsache weist darauf hin, dass das Eindringen von Lehnwörtern aus dem Arabischen, d. h. einer Sprache anderer Struktur, die Struktur des Neupersischen nicht geändert hat. Die Anwendung der Lehnwörter in bezug auf die Wortlänge nach der Silbenzahl unterliegt den strukturellen Gesetzmässigkeiten des Neupersischen, das in dieser Hinsicht eine unmittelbare Kontinuation des Mittelpersischen ist.

Die Vergrößerung der Wirksamkeit des Codes der persischen Sprache bedeutet, dass sie im Laufe ihrer Entwicklung immer besser an die Forderung 'je öfter ein Symbol vorkommt, desto kürzer soll es sein' angepasst wird. Diese Anpassung erfolgt jedoch nicht durch einfache Verkürzung der Wörter. Dagegen spricht die Tatsache, dass die grösstmögliche Wortlänge, sowohl nach der Silben- als auch nach der Phonemzahl, in allen drei untersuchten Sprachen gleich bleibt: 6 Silben und 14—16 Phoneme.

¹ R. Koppe, *Statistik und Semantik der arabischen Lehnwörter in der Sprache 'Alavi's*, Wiss. Z. Humboldt-Univ. Berlin, Ges.-Sprachw. R. IX (1959/60), S. 585—619; W. Skalmowski, *Ein Beitrag zur Statistik der arabischen Lehnwörter im Neupersischen*, Folia Orientalia, III (1961), S. 171ff.

² W. Fucks, *Mathematische Analyse von Sprachelementen, Sprachstil und Sprachen*, Köln u. Opladen 1955, S. 84.

Die Vergrößerung der Wirksamkeit kann auf zweierlei Weise erfolgen: 1) durch die Verkürzung der längeren Wörter, die häufig vorkommen und 2) durch Ersetzen der längeren Wörter durch andere, kürzere.

In der persischen Sprache hat man mit den beiden Prozessen zu tun. Die Verkürzung der Wörter tritt besonders deutlich beim Übergang vom Altpersischen zum Mittelpersischen zutage (Abfall der auslautenden Silbe). Dieser Prozess wird deutlich durch die Charakteristika M_s und M_{ph} widergespiegelt. Beim Übergang vom Mittelpersischen zum Neupersischen ist eine direkte Verkürzung der Wörter weniger augenscheinlich. Die Tatsache jedoch, dass der neupersische Wortschatz ca. 50% arabischer Lehnwörter enthält, die Wirksamkeitsindices des Neupersischen aber trotzdem weiter steigen, weist darauf hin, dass manche iranische Wörter durch arabische Lehnwörter ersetzt werden und zwar in Übereinstimmung mit dem Kriterium der Wirksamkeit.

Es wurde bereits darauf hingewiesen, dass die Charakteristika der Worttypen im Arabischen stark von den entsprechenden Charakteristika des Mittel- und Neupersischen abweichen. Die arabischen Lehnwörter werden jedoch im Neupersischen so angewandt, dass die Struktur des Neupersischen und die Kontinuität der persischen Sprachentwicklung durch sie nicht beeinträchtigt werden. Es scheint also, dass die Tendenz zur Vergrößerung der Codewirksamkeit, die im Persischen deutlich zu sehen ist, einen Einfluss auf die Anwendung der Wörter ausübt.

Die Tendenz zur Vergrößerung der Codewirksamkeit und ihr Einfluss auf die Anwendung und Veränderungen der Wörter werden auch in anderen Sprachen beobachtet. Die auftretenden Veränderungen werden meistens durch die Mitglieder der Sprachgemeinschaft als „Sprachverderbnis“ empfunden. Eine gute Darstellung der Wirkung dieser Tendenz im Deutschen bietet die Abhandlung von A. Schopenhauer u. d. T. „Über die, seit einigen Jahren, methodisch betriebene Verhunzung der deutschen Sprache“¹. Die Hauptzüge dieses Prozesses, die Schopenhauer vom ästhetischen Standpunkt aus scharf verurteilt, werden von

¹ Schopenhauer's handschriftlicher Nachlass, II Bd., hg. E. Griesbach, Ph. Reclam, Leipzig o. D.; angeführte Stellen S. 118 und 125f.

ihm folgendermassen geschildert: „Eine fixe Idee hat sich aller deutschen Schriftsteller und Schreiber jeder Art, vielleicht mit wenigen, mir nicht bekannten Ausnahmen, bemächtigt: sie wollen die deutsche Sprache zusammenziehen, sie abkürzen, sie kompakter, knapper machen. Zu diesem Ende ist ihr oberster Grundsatz, überall das kürzere Wort dem gehörigen oder passenden vorzuziehen. Es wird bald auf Kosten der Grammatik, bald auf Kosten des Sinnes, dann also lexikalisch, endlich und wenigstens auf Kosten des Wohlklangs durchgesetzt, und zwar so, dass sie sich Gewaltthätigkeiten jeder Art gegen die Sprache erlauben: sie muss biegen oder brechen“. „... Der schmutzigste Buchstabengeist beherrscht sie (d. h. die Sprachverhunzer). Ihr leitender Grundsatz ist: nicht das richtige Wort, sondern das kürzere, wenn es nur so *à peu près* die Sache bezeichnet: dem Leser bleibt überlassen, unsere Meinung zu errathen“.

Schopenhauer schildert hier auf anschauliche Weise den Verlauf des Eindämmens des „sprachlichen Überflusses“, das durch die Redundanz verschiedener Sprachschichten ohne Unterbrechung der Verständigung möglich ist.

In der vorliegenden Untersuchung wurde nur die Beziehung zwischen der Information und der signifiant-Seite des sprachlichen Zeichens berücksichtigt. Da jedoch die selektive Information mit dem sprachlichen Wert gleichgesetzt werden kann, und dieser den Ausdrucksplan mit dem Inhaltsplan des Zeichens verbindet, kann eine Wechselwirkung der beiden Seiten des sprachlichen Zeichens auf dessen Informationsinhalt und damit auf die Wirksamkeit des Codes angenommen werden. Damit soll gesagt werden, dass die beobachtete Tendenz zur Vergrößerung der Wirksamkeit des Codes, die man als Wirkung des Prinzips des geringsten Kraftaufwandes in der Sprache im Sinne von Zipf¹ interpretieren könnte, nicht nur eine Begleiterscheinung eines Bedeutungswandels ist, der allein den Übergang von einem Sprachtypus zum anderen — z. B. vom synthetischen zum analytischen, wie der Übergang vom Altpersischen zum Mittelpersischen — verursachen kann, sondern als mitwirkender Faktor der Sprach-

¹ G. K. Zipf, *Human Behavior and the Principle of Least Effort*, Cambridge, Mass., 1949.

entwicklung angenommen werden muss. Darauf weist u. a. die Tatsache hin, dass die Werte der Wirksamkeitsindices beim Übergang vom Mittelpersischen zum Neupersischen steigen, obwohl diese beiden Sprachen analytisch sind und die grammatische Rolle, ihre „Belastung mit grammatischen Funktionen“ in beiden Sprachen gleich ist.

Die Untersuchung der Entwicklung der persischen Sprache zeigt also, dass diese Entwicklung nicht vollkommen indeterminiert ist, sondern dass die Beziehung zwischen dem sprachlichen Wert des Zeichens und seiner signifiant-Seite einem ständig wirkenden, panchronen Faktor unterliegt.

FRANCISZEK MACHALSKI

VAḤĪD DASTGARDĪ AND HIS „ARMAĠĀN“

The name of Vaḥīd Dastgardī, one of the most enlightened and progressive representatives of the Persian intellectuals of the present generation glitters with its specific blaze in the firmament of the contemporary Iranian literature and culture. Publicist, literary man and poet in one person, indefatigable worker in the field of the native literature, he passed to the history as one of its creators and renovators.

Hasan with his literary surname Vaḥīd came forth out of the Persian people. He was born in 1298 h./1880 of our era in the village Dastgard situated at a distance of 1 parasang in the south from Isfahan, not far away from the suburb Ġolfā. His father, Abo'l-Kāsem, was a poor farmer, but he managed to ensure his son a convenient education. Hasan, at the age of fifteen, went to Isfahan and learned there, for ten years, the Arabic language, literature and philosophy.

When the time of struggle for the constitution came, Hasan having taken the name of Dastgardī left his school and joined the ranks of constitutionalists. He offered all his literary gift and inborn abilities to the great cause of freedom and democratic liberties. The Isfahan periodicals: „Parvāneh“ (Butterfly), „Zājan-deh-rūd“¹, „Mofatteš-e Īrān“ (Iranian Inspector) swarmed with his poems and articles on social and political subjects.

At the onset of the First World War Vaḥīd started to publish and edit periodical of his own entitled „Derafš-e Kāviyān“ (Banner of Kaveh). His political sympathies were then on the side of the allied Germany, Austria and Turkey, to which

¹ The name of the river flowing through Isfahan.