

NORTH-WESTERN EUROPEAN
LANGUAGE EVOLUTION
Supplement vol. 23



North-Western European Language Evolution (NOWELE) is a scholarly journal which publishes articles dealing with all aspects of the histories of – and with intra- and extralinguistic factors contributing to change and variation within – Icelandic, Faroese, Norwegian, Swedish, Danish, Frisian, Dutch, German, English, Gothic and the Early Runic language. *NOWELE* is edited by Michael Barnes (London), Rolf H. Bremmer Jr. (Leiden), Gotthard Lerchner (Leipzig) and Hans F. Nielsen (Odense).

Beiträge zur Morphologie

Germanisch, Baltisch, Ostseefinnisch

Herausgegeben
von

Hans Fix

UNIVERSITY PRESS OF SOUTHERN
DENMARK
2007

ISBN 978 87 7674 249 2 (Pb ; alk. paper)

ISBN 978 90 272 7272 0 (Eb)

© 2012 – John Benjamins B.V.

Published 2007 by University Press of Southern Denmark

No part of this book may be reproduced in any form, by print, photoprint, microfilm, or any other means, without written permission from the publisher.

John Benjamins Publishing Co. · P.O. Box 36224 · 1020 ME Amsterdam · The Netherlands
John Benjamins North America · P.O. Box 27519 · Philadelphia PA 19118-0519 · USA

INHALTSVERZEICHNIS

VORWORT.....	vii
ENDRES, Norbert: Algorithmische Sprachwissenschaft - rechnergestützte Darstellung von Grammatik und Modellierung von Phänomenen des Lautwandels	1
HEIDERMANNS, Frank: Zum Präteritum der starken Verben im Gotischen	57
QUAK, Arend: Zum <i>-in</i> -Suffix in den Malbergischen Glossen der <i>Lex Salica</i>	69
DIETZ, Klaus: Denominale Abstraktbildungen des Altenglischen: die Wortbildung der Abstrakta auf <i>-dōm</i> , <i>-hād</i> , <i>-lāc</i> , <i>-rāden</i> , <i>-sceaft</i> , <i>-stæf</i> und <i>-wist</i> und ihrer Entsprechungen im Althochdeutschen und im Altnordischen	97
KORNEXL, Lucia: Zur Rolle der Morphologie in den zweisprachigen Wörterbüchern des Spätmittelenglischen.....	173
KLEIN, Thomas: Von der semantischen zur morphologischen Steuerung. Zum Wandel der Adjektivdeklinaton althochdeutscher und mittelhochdeutscher Zeit.....	193
MÖHN, Dieter & SCHRÖDER, Ingrid: Sprachbedarf und Lexembildung am Beispiel der Grammatik des Mittelniederdeutschen	227
KROGH, Steffen: Zur Diachronie der nominalen Pluralbildung im Ostjiddischen	259
DE LEEUW VAN WEENEN, Andrea: Morphologische Entwicklung im isländischen Pronominalsystem	287
FIX, Hans: Flexionsmorphologische Varianz in der altisländischen <i>Elucidarius</i> -Tradition	321
SCHABALIN, Andreas: Das Greifswalder altnordische morphologische Wörterbuch	335
SKRZYPEK, Dominika: Reductions in Morphological Systems: The Case of Old Swedish.....	351
GÖTZSCHE, Hans: <i>Die Bibel Christians III.</i> : Zwischen Morphologie und Topologie – eine vorläufige Hypothese.....	361
FECHT, Rainer: Lit. <i>pláuti</i> : aksl. <i>pluti</i> - eine Frage der Morphonologie.....	383
RANGE, Jochen D.: Zur synchronen Interpretation flexionsmorphologischer Varianz in den Evangelien und der Apostelgeschichte der altlitauischen Bibelübersetzung 1579-1580 von Johannes Bretke.....	395
GRÜNTAL, Riho: Morphological Change and the Influence of Language Contacts in Estonian	403
LAAKSO, Johanna: Transitivity und Intransitivity in der Wortbildung: Estnisch und Deutsch.....	433
PANTERMÖLLER, Marko: Reanalyse, Reaktionsanalogie oder Formensynkretismus? Reaktionswandel in Abessivität ausdrückenden Adpositionalphrasen am Beispiel finnischer Verordnungs- und Erlaßtexte zwischen 1584 und 1810.....	449

VORWORT

Eine interdisziplinäre Tagung *Morphologische Probleme in den Sprachen der Ostseeanrainer* war das Vortragsforum für die hier zusammengestellten, gegenüber der Vortragsfassung zum Teil stark erweiterten Aufsätze. Fast alle hier versammelten Autoren waren vom 14. bis 17. September 2005 der Einladung an das Alfred-Krupp-Wissenschaftskolleg Greifswald gefolgt, um nicht nur einander, sondern auch eine kleine interessierte Öffentlichkeit über ihre morphologischen Forschungen zu unterrichten und zugleich Einblick zu nehmen in das am Kolleg angesiedelte Forschungsprojekt zur Morphologie des Altnordischen.

Der Ostseeraum, der das besondere Interesse der Greifswalder Universität genießt, wurde dabei weit ausgelegt; neben den Anrainersprachfamilien im engeren Sinn Baltisch, Slawisch und Ostseefinnisch sollten möglichst alle altgermanischen Sprachen bis zur Reformation berücksichtigt werden, selbst die ostseefernen. Da es uns trotz wiederholter Bemühung nicht geglückt war, zum Tagungszeitpunkt Referenten zu gewinnen, sind die slawischen Sprachen hier leider nur peripher vertreten.

Damit gliedert sich der vorliegende Band in vier Themenbereiche:

1. ein computerlinguistischer Beitrag (Endres) appelliert programmatisch an die Sprachwissenschaft, ihre Aussagen über Formenbildung und Sprachwandel präzise zu formulieren, damit sie auf Rechner modellierbar und so nachprüfbar werden,
2. elf Beiträge stellen das Ostgermanische (Heidermanns), das Westgermanische (Quak, Dietz, Kornexl, Klein, Möhn/Schröder, Krogh) und das Nordgermanische (de Leeuw van Weenen, Fix, Schabalin, Skrzybek, Götzsche) in den Mittelpunkt,
3. zwei Beiträge beschäftigen sich mit dem Baltoslawischen (Fecht) und dem Litauischen (Range),
4. das Ostseefinnische ist mit drei Beiträgen zum Estnischen (Grünthal, Laakso) und zum Finnischen (Pantermöller) vertreten.

Als Herausgeber bleibt mir zum Schluss die angenehme Pflicht, dem Alfred-Krupp-Wissenschaftskolleg Greifswald für die Gastfreundschaft im Kolleg und für die Finanzierung sowohl der Tagung als auch der Drucklegung dieses Bandes herzlich zu danken. Mein Dank geht natürlich an die Autorinnen und Autoren, für ihre Bereitschaft mitzuarbeiten und für ihre Geduld, die bei Sammelbänden immer wieder gefordert ist. Schließlich gebührt ein besonderes Dankeschön meiner Mitarbeiterin, Dr. Birgit Hoffmann, die mit großer Sorgfalt die Texte des Bandes einheitlich formatierte und dabei noch manche „Kleinigkeit“ bemerkte und ausmerzte, die allen zuvor entgangen war.

Greifswald, im Mai 2007
Hans Fix

Norbert Endres

ALGORITHMISCHE SPRACHWISSENSCHAFT – RECHNERGESTÜTZTE DARSTELLUNG VON GRAMMATIK UND MODELLIERUNG VON PHÄNOMENEN DES LAUTWANDELS

Wenn algorithmische Sprachwissenschaft als Titel für diese Arbeit gewählt wurde, so ist dies nicht geschehen, um der wachsenden Zahl als wissenschaftlich ausgerufenen Disziplinen eine weitere hinzuzufügen, sondern vielmehr in der Absicht, durch die getroffene Wortwahl darauf einzuwirken, daß die nun bereits seit mehreren Jahrzehnten zur Verfügung stehenden und sich anbietenden, aber ungenutzt gebliebenen Möglichkeiten für die historische und vergleichende Sprachwissenschaft auch tatsächlich wahrgenommen werden, wie dies in einem geneigteren wissenschaftspolitischen Umfeld sicherlich längst verwirklicht worden wäre. Die technischen Voraussetzungen zumindest sind verfügbar, um umfangreichen Regelwerken große Datenmengen zur Überprüfung zuzuführen und damit in Anlehnung an eine naturwissenschaftliche Vorgehensweise die nur aus kleineren Auszügen sprachlichen Materials abstrahierten Regeln an hinreichend großen und die Gesamtheit der sprachlichen Realität repräsentierenden Materialproben zu verifizieren.

1. *Eine Konzeption zur diachronen Sprachwissenschaft*

Die nachstehenden Abschnitte geben die Sichtweise eines theoretischen Mathematikers auf die Forschungsgegenstände der historischen und vergleichenden Sprachwissenschaft wieder, wobei ausdrücklich historische und vergleichende Sprachwissenschaft nicht alleine als Synonym für Indogermanistik, sondern im allgemeinsten Wortsinne verstanden werden soll.

Es wird in diesem Aufsatz die auch von manchen sprachwissenschaftlichen Kreisen vertretene Ansicht geteilt, daß die historische und vergleichende Sprachwissenschaft weit hinter den seit mehreren Jahrzehnten durch elektronische Datenverarbeitung und die Theorie formaler Sprachen eröffneten Möglichkeiten geblieben ist und das hohe ihr innewohnende Potential an Aussage-

kraft nicht annähernd ausgeschöpft hat. Wesentliche Schritte hin zur Ausschöpfung dieses Potentials bestehen in der klaren Abgrenzung konkurrierender Denkmodelle voneinander, der Auflösung von Implizitheiten und der Disambiguierung von Regeln zum Lautwandel und zur Formenbildung (vgl. Rix et al. 1998:4) sowie der exakten Formulierung von Analogien. Hierbei kann nicht darauf gehofft werden, daß etwa Vertreter der Computerlinguistik, die an sprachhistorischen Prozessen und der damit verbundenen Detailarbeit bestenfalls marginal interessiert sind, von sich aus entsprechende Anstrengungen unternehmen würden, die vorhandenen Lücken zu schließen. Umgekehrt ist es von geringem Wert, wenn Regeln zum Lautwandel in der herkömmlichen, unzureichend präzisen Form an die Darstellungsweisen bestehender formaler Ansätze, wie generative Grammatiken, Optimalitätstheorie, oder was sonst gerade berechtigt oder unberechtigt den Anspruch auf Fortschrittlichkeit erhebt, nur angepaßt werden; dies führt lediglich zu Paraphrasierung ohne Erkenntnisgewinn.

Grundsätzlich scheint es geboten, philologisch-sprachwissenschaftliche und außerphilologische, formal-abstrakte Denkmuster zu kombinieren, und es ist angeraten, das bisher nur vage umschriebene „System der Sprache“ in seinen formalisierbaren Teilen in Datenstrukturen abzubilden und damit in Wörterbüchern und Grammatiken verzeichnete Information explizit, greifbar und abrufbar zu machen. Mit den so gewonnenen Mitteln können Modelle, und vielleicht schrittweise alle miteinander konkurrierenden Modelle, zur diachronen Sprachentwicklung aufgestellt werden, die mit der durch Zeugnisse vermittelten Realität zumindest nicht in Widerspruch stehen, und man erreicht es, Etymologien zu verifizieren, soweit es nur die sich in der Graphie widerspiegelnden phonetischen und morphologischen Kriterien im Sinne Abaevs (1980: 36) betrifft, d.h. der Feststellung gesetzmäßiger Entsprechungen von Lautungen und Wortbildungsformantien. Die Verifikation darüber hinausgehender Aspekte von Etymologien, beispielsweise der Bedeutungsveränderung oder kulturhistorischer Angaben jenseits von Sprachkontaktphänomenen, liegt gegenwärtig außerhalb der Reichweite einer solchen Abbildung in Datenstrukturen.

In größtmöglicher Abstraktion formuliert besteht das hier vorgetragene Interesse in der Realisierung einer Konzeption, mit der unter Verwendung von formalen Mitteln und daraus hervorgehenden, auf Rechnern implementierten Algorithmen für Dialekte einer Sprachfamilie die Entwicklung für alle denkbaren Wortformen, von einer rekonstruierten oder belegten Grundsprache aus-

gehend, für ein geeignet zu wählendes Zeitintervall modelliert wird, wobei sowohl Lautgesetze als auch Analogien, Wechsel der Flexionsklasse und Synkretismen Berücksichtigung finden.

Hierzu wird ein die Komponenten (1) und (2) umfassendes iteratives Vorgehen vorgeschlagen, beginnend mit separat aufzustellenden Ausgangspunkten für (1) und (2), und einer nach Anpassung der Schreibung an die betrachteten Sprachstufen und den erforderlichen Aufstockungen des Wortformeninventars früh statisch werdenden Komponente (3) zur Kontrolle.

- a) Ansetzen von Flexion und Lexik der betrachteten Grundsprache und Überprüfung ihrer Rekonstrukte oder Belege daraufhin, ob einander entsprechende Wortformen der Dialekte aus ein und demselben Rekonstrukt, Transponat oder Beleg der Grundsprache unter Berücksichtigung von (2) hergeleitet werden können.
- b) Formalisierung und Verifikation der in historischen Grammatiken nur in paraphrasierter oder exemplarischer Form gegebenen Transformationsregeln (Lautgesetze und Analogien); Abgleich sowohl der genannten Transformationsregeln als auch der relativen Chronologien der betrachteten Dialekte gegeneinander unter Verwendung von (1).
- c) Wortformenkatalog in normalisierter Schreibung für die Dialekte für unterschiedliche Schnitte innerhalb des gewählten Zeitintervalls als Zielpunkte für die aus der kombinierten Anwendung von (1) und (2) gewonnenen Ergebnisse.

Mit einem derartigen Vorgehen ermöglicht man eine wechselseitige Korrektur der Angaben in für unterschiedliche Sprachfamilien oder Dialekte oftmals voneinander unabhängig erstellten Grammatiken und etymologischen Wörterbüchern.

Eine solche Konzeption ist universell und somit auf beliebige Sprachfamilien anwendbar, sie ist ein kanonisches Mittel, der vergleichenden Sprachwissenschaft zu einer zusätzlichen, ihr adäquaten Basis von einer Art zu verhelfen, wie sie für die Naturwissenschaften selbstverständlich, aber auch der Sprachwissenschaft keineswegs wesensfremd ist. Für den Bereich der Indogermania etwa würde man auf diese Weise eine Verifikation und Verbesserung der bisher niemals mit formalen Mitteln nachvollzogenen Aussagen der Junggrammatiker und ihrer Nachfolger erreichen, wenn man als Grundsprache Urindogermanisch wählen würde und als Dialekte die Urformen der indogermanischen Sprachzweige.

Damit kann auch der übliche sprachwissenschaftliche Beweisbegriff auf eine andere Ebene verlagert werden. Traditionell besteht dieser Beweisbegriff lediglich in der Gestalt eines auf vergleichsweise wenigen Exemplaren fußenden und Verallgemeinerung beanspruchenden Indizienbeweises, bei dem oftmals die Prämissen, manchmal sogar die Bewertung von Bedingungen als notwendige oder hinreichende, vage bleiben.

Die gerade vorgestellte Konzeption wurde inspiriert und maßgeblich beeinflusst durch das erstmalig in Fix & Poschenrieder (1998:30-40) vorgestellte RWH-Konzept eines Rechnergestützten Wörterbuches mit Handschriftenanbindung, bei dem für eine Einzelsprache und ihre Vorgängersprache die Stimmigkeit zwischen den Angaben in den Wörterbüchern und historischen Grammatiken einerseits und der durch Überlieferung in Handschriften tradierten sprachlichen Realität andererseits mit datentechnischen Hilfsmitteln überprüft wird. Es soll insofern über das RWH-Konzept hinausgegangen werden, als gleich ein ganzes Bündel von verwandten Sprachen untersucht und als Datengrundlage auch ausschließlich virtuelles, also rekonstruiertes Wortformenmaterial zugelassen wird.

Aus dieser Konzeption erwachsende konkrete mittelfristige Forschungsvorhaben wären,

- die Entwicklung zunächst einiger als enger zusammengehörig erachteter Sprachzweige des Indogermanischen, wie Indisch und Iranisch oder Baltisch und Slawisch, aus der gemeinsamen Grundsprache rechnergestützt zu modellieren,
- das im Lexikon indogermanischer Verben (LIV, Rix et al. 1998) dargebotene Material auszuwerten, indem von der Kombination aus indogermanischen Verbalwurzeln und der vorgenommenen Typisierung in indogermanische Flexionsklassen ausgehend an den gegebenen Beispielen von Primärverben die Entwicklung in die Einzelsprachen rechnergestützt nachvollzogen wird,

um die Grundlage dafür zu schaffen, als Fernziele schließlich

- die Sprachentwicklung über die gesamte Indogermania hinweg geschlossen und einheitlich darzustellen,
- durch eine vollständige Zusammenstellung der Regeln des Lautwandels für die indogermanischen Sprachen eine Ausgangsbasis für eine Typologie des Lautwandels zu erhalten, die auch die Erforschung der nicht durch Belege dokumentierten Entwicklung nichtindogermanischer Sprachen erleichtern würde (Tichy 2000:24).

Im weiteren Verlauf der Arbeit wird der Leistungsumfang einiger umfangreicherer Algorithmen vorgestellt und die aus ihnen gewonnenen Ergebnisse kommentiert. Es erfolgt allerdings weder eine Beschreibung der Algorithmen selbst, noch eine eingehendere Erörterung der Transformationsregeln, da dies trotz erster Verifikationsschritte wegen noch ausstehender Untersuchungen verfrüht wäre. Die von den Algorithmen behandelten Fragestellungen sind

- Formenerzeugung indogermanischer Verben
- Modellierung des Lautwandels von einer Grundsprache in einen Nachfolgedialekt
- Formenerzeugung altisländischer schwacher Verben

Nur aus oberflächlicher Sicht dienen die Programme zur Formenerzeugung lediglich der rechnergestützten Darstellung von Auszügen einzelsprachlicher grammatischer Erscheinungen, doch sind auch hier Routinen zur Behandlung aller zu berücksichtigenden diachronen Vorgänge implementiert.

Welche der von der Computerlinguistik entwickelten Formalismen geeignet sind, Sprachwandel adäquat modellieren zu können, kann vorab, ohne umfassende Kenntnis aller auftretenden Phänomene, nicht befriedigend geklärt werden (Carstensen 2001:20). Die hier vorzustellenden Algorithmen beschränken sich auf algebraisch formulierbare Termersetzungsregeln und bewegen sich im Rahmen der zweiwertigen Logik.

Mit dem in den Kapiteln 3 bis 5 präsentierten Material wird gezeigt, daß das den Sprachwissenschaften innewohnende, bisher unausgeschöpft gebliebene Potential an Aussagekraft erschließbar ist und demnach die Erfolgsaussichten für die vorgestellte Konzeption und die aus ihr ableitbaren Forschungsvorhaben hoch sind.

2. *Sprachwissenschaftliche Literatur als Grundlage zur Programmierung*

Zunächst ist, wenigstens punktuell, die Zuverlässigkeit sprachwissenschaftlicher Angaben in der Fachliteratur und ihre Brauchbarkeit für die formale Umsetzung, wie sie für die Programmierung erfolgen muß, zu beleuchten. Leider entsteht bei der Lektüre der Eindruck, daß ein Mangel an Explizitheit und an erreichbarer Präzision nicht wenigen als Kavaliersdelikt zu gelten scheint. Ein eingehenderes Studium etymologischer Wörterbücher und historischer Grammatiken lehrt schnell, daß man nicht selten und nicht nur in Detailfragen im Stich gelassen wird, daß Autoren vielfach die Prämissen nicht oder nur unvoll-

ständig offenlegen, unter denen sie ihre Aussagen treffen, daß Informationen häufig nur exemplarisch, fragmentarisch oder implizit gegeben werden oder ohne vom Autor stumm vorausgesetztes Wissen oder intimerer Kenntnis der Schule, aus der der Autor hervorgegangen ist, nicht verständlich sind, daß innerhalb weniger Seiten widersprüchliche Angaben gemacht werden, oder daß zeitlich inhomogenes Material nebeneinandergestellt wird. Selbst auf modernere Arbeiten, die unter dem Einfluß oder aus der Perspektive von Theorien zur formalen Behandlung sprachlicher Phänomene verfaßt wurden, trifft dies noch zu, wodurch sich bei unzureichender Analyse höchst unerfreuliche Mißerfolge in der Programmierung einstellen können. Es folgen einige Beispiele, die als symptomatisch für die oben angemahnten Mängel gelten können und teilweise für den weiteren Verlauf dieses Aufsatzes von Bedeutung sind.

Eine der wenigen Arbeiten, in der eine Formalisierung der das Altisländische betreffenden Regeln zum Lautwandel versucht wird, ist Catheys Dissertation „A Relative Chronology of Old Icelandic Phonology Based on Distinctive Feature Analysis“ (Cathey 1967), die unter dem Einfluß der Sichtweisen von Chomsky, Jakobson und Halle steht. Besprochen werden insgesamt 106 Transformationsregeln des Lautwandels vom Urgermanischen hin zum Altisländischen, deren Wirkung an 406 verschiedenen Wortformen demonstriert wird. Sehr anerkennenswert ist die überall spürbare Absicht, die nur in Prosaform ausgeführten oder lediglich durch einige Beispiele gestützten Aussagen zum Lautwandel, meist aus Heusler (1964) oder Noreen (1923), zu präzisieren und in eine zeitliche Abfolge zu bringen.

Die pseudoformale Gestalt allerdings, in der die Lautwandelregeln niedergelegt sind, vermittelt eine trügerische Sicherheit. Zur korrekten Formalisierung oder Programmierung können sie nicht ohne weiteres direkt aus Catheys Arbeit übernommen werden, da zuviel Interpretationsspielraum bleibt, den man aus den dargebotenen Wortbeispielen zusammen mit den gleichfalls nicht sauberen Merkmalsdarstellungen einzugrenzen versuchen muß. Exemplarisch sei dies anhand von Cathey (1967:8) illustriert.

Rule 1C: Short i is lost in the third syllable

Zeichenkettendarstellung unter Verwendung segmentaler Symbole

$$i \quad \rightarrow \quad \emptyset \quad / \quad [(V)] \quad \begin{matrix} [(')] \\ [V] \end{matrix} \quad [C] \quad ___\quad (C) \quad [\#]$$

Merkmalsdarstellung oder merkmalsbasierte Darstellung

[+son]	→ [-son]	/ [(-cons)]	[-cons]	[-voc]	[-cons]	([+fea])	[-cons]
[+diff]	[-diff]	[(+son)]	[+son]		[-grav]		[-son]
[+short]	[-short]		[-accI]		_____		[-tens]
[+fea]	[-fea]						

Formal unhaltbar sind die Merkmalsdarstellungen zu $\emptyset$ und der mit $\#$ bezeichneten Wortgrenze sowie die unterschiedlichen Auslegungen sowohl für das Minuszeichen als auch für das Pluszeichen, und bei einem weiteren Blick auf die Merkmalsdarstellung erhebt sich die Frage, ob in der Zeichenkettendarstellung die beiden C jeweils stellvertretend für dieselbe Gruppe von Lauten stehen; die Angabe $([+fea])$ in der Merkmalsdarstellung zu (C) ist da wenig erhellend, wie es überhaupt der Gebrauch von weiter nicht spezifizierendem fea mit wechselnden Vorzeichen auch an anderen Stellen der Merkmalsdarstellung ist.

Schon diese keineswegs vollständige Analyse zeigt, daß der Versuch, eine Transformationsregel schlüssig niederzulegen, hier mißlungen ist. Zwar erinnert Regel 1C durch ihre Gestalt an eine Termersetzung, bei anderer Sichtweise an eine Aussage, aber nur unter Zuhilfenahme von schriftlich nirgends fixierten Zusatzannahmen kann 1C in eine widerspruchsfreie Form gebracht werden, in der die den segmentalen Symbolen und den Merkmalen beizumessenden Bedeutungen miteinander in Einklang stehen.

Das Studium der Beispiele zu Regel 1C in Cathey legt immerhin nahe, die Zugehörigkeit von j zu $[C]$ der Zeichenkettendarstellung positiv zu entscheiden, denn

Rule 1C is a familiar PGmc rule which says that such early forms as **berizi* and **zastijiz* became respectively **biriz* [sic!] and **zastijz* (*zastīz*).

Prinzipiell und wenig befriedigend äußert sich Beekes zu derartigen Formaldarstellungen (Beekes 1995:60).

The formulae are difficult to draw (and expensive to reproduce), and their meaning is not always immediately apparent [...]. A simpler procedure is simply to forget about the formula and to describe the process of sound change verbally. It is, by the way, quite possible after a little practice to understand what happened phonetically without having to be told.

Die von Cathey vorgelegte Aufstellung von Transformationsregeln ist sehr breit angelegt, ihre Ausarbeitung hat erkennbar viel Mühe bereitet und verdient zweifellos weitere Beachtung. Trotz einiger formaler Mängel, und obwohl sich im Sinne der obigen Kritik an sprachwissenschaftlicher Literatur noch viele weitere Unzulänglichkeiten bei Cathey entdecken lassen, erwies sich seine Dissertation als insgesamt so vertrauenswürdig, daß sie für zwei der hier vorzustellenden Programmiervorhaben, eine Modellierung des Lautwandels vom Urgermanischen zum Altisländischen und eine Wortformenerzeugung für altisländische schwache verbale Simplizia, als eine wesentliche Grundlage dient, wobei hinzuzufügen ist, daß sowohl die beiden Programmiervorhaben als auch die hierfür aus Cathey abgeleiteten Daten voneinander unabhängig erstellt sind.

Es sind in der sprachwissenschaftlichen Literatur jedoch nicht nur solche Ungereimtheiten feststellbar, deren Aufdeckung und, sofern dies möglich ist, Richtigstellung oder Disambiguierung einen etwas erhöhten analytischen Aufwand erfordern; oftmals sind es schlichte und vermeidbare Disziplinlosigkeiten oder Läßlichkeiten der Autoren, die einer Auswertung der dargebotenen Angaben entgegenstehen oder sie wenigstens maßgeblich erschweren. Zur Bildung des indogermanischen Imperfekts lehrt Beekes (1995:226), daß das der Verbalwurzel stets voranzustellende, Augment genannte Formans *h₂e* immer betont ist (*'the augment took the stress'*), um selbst 31 Seiten später ohne Nennung von Gründen Imperfektformen zu rekonstruieren, deren Betonung ganz offensichtlich auf die Endung zu liegen kommen soll. Szemerényi (1990) führt im Kapitel zur Phonologie zwar für nichtsyllabische *u* und *i* die Symbole *w* und *y*, sowie für syllabische *m*, *n*, *r*, *l* die Symbole *m*, *n*, *r*, *l* ein, gebraucht im weiteren Verlauf seines Buches aber dann ohne nähere Erläuterung diese Symbole oder ersetzt sie durch *u*, *i*, *m*, *n*, *r*, *l*; einer gewissen Beliebigkeit folgt die Markierung der Betonung.

Stellenweise könnte man zu dem Eindruck gelangen, daß ein höherer Grad an Unschärfe, als der Untersuchungsgegenstand selbst es erfordert, zum Pro-

gramm erhoben wird und eine nicht unerhebliche Furcht davor besteht, ein Optimum an Präzision erreichen zu wollen, wie sich auch manchen Passagen aus Szemerényis „Einführung in die vergleichende Sprachwissenschaft“ entnehmen läßt, wo von *‘Edgertons vielleicht allzu präzisen Formeln’* die Rede ist (Szemerényi 1990:115), die zu berichtigen seien, oder argumentiert wird, daß eine *‘... Erklärung, obwohl sie sehr geistreich ist, dennoch nicht befriedigt und keine wirkliche Lösung bietet’*, was aber *‘... nicht nur an der allzu mathematischen Einstellung [liegt]’* (Szemerényi 1990:120). Meines Erachtens ist diese Furcht völlig unangebracht. Natürlich kann es um die Richtigkeit von Edgertons Regeln schlecht bestellt sein (Sihler 2006), wenn sie in nicht materialgerechter Weise überspezifiziert sind; die Gründe hierfür aber sind weder in der formalen Darstellung oder in allzu präzisen Formeln noch in einer mathematischen oder gar allzu mathematischen Einstellung zu suchen. Eine in sich schlüssige Darstellung sprachlicher Phänomene muß formal korrekt sein, womit von einer durch ein Übermaß an Präzision ausgehenden Gefahr keine Rede sein kann; wenn eine Darstellung sprachlicher Phänomene scheitert, sind Ursachen hierfür im Mangel an Präzision und in Unvollständigkeit der Beschreibung zu suchen.

3. *Indogermanische Verbalformenwerkbank*

3.1. *Aufbau und Zielsetzung*

Mit der indogermanischen Verbalformenwerkbank wird die Absicht verfolgt, die Bildung finiter indogermanischer Verbalformen aus den ihnen zugrundeliegenden Morphemen transparent zu machen. Es wurden möglichst viele theoretische Sichtweisen aus der Standardliteratur als Grundeinstellungen in das Programm übernommen, sofern diese Sichtweisen nur hinreichend detailliert ausformuliert und aufeinander abstimmbare waren. Der Benutzer kann sehr weitgehend Einfluß auf die Gestalt der Formantien nehmen; eine Einflußnahme auch auf die Zeichenkettenmanipulationen, die hier nur im Rahmen der Syllabierungsregeln zu berücksichtigen sind, ist allerdings vorerst nicht gegeben.

Das Programm befindet sich in einem zwar fortgeschrittenen, trotzdem aber noch prototypisch zu nennenden Entwicklungszustand. Es ist erhebliche weitere Verifikations- und Verbesserungsarbeit zu leisten, wenn man es dem Benut-

zer ermöglichen möchte, ein Maximum an eigenen Vorstellungen hinsichtlich der Formenbildung zuverlässig modellieren zu lassen.

Zur Wiedergabe von Wortformen wird, abgestimmt auf die derzeit einbezogenen Theorien, das nachstehend spezifizierte Alphabet benutzt. Hierbei ist Akut als Diakritikum äquivalent zu Betonung, Überstrich als Diakritikum ist äquivalent zu Länge. Ein untergestellter Kreis impliziert das Vorliegen eines syllabischen Elements, ein untergestellter Halbkreis das eines nichtsyllabischen Elements.

Tenues	$p, t, \hat{k}, k, k^u$
Tenues aspiratae	$p^h, t^h, \hat{k}^h, k^h, k^{uh}$
Mediae	$b, d, \hat{g}, g, g^u$
Mediae aspiratae	$b^h, d^h, \hat{g}^h, g^h, g^{uh}$
Sibilanten	s
Laryngale, nichtsyllabisch	h_1, h_2, h_3
Laryngale, syllabisch	$\underset{\circ}{h}_1, \underset{\circ}{h}_2, \underset{\circ}{h}_3$
Nasale und Liquiden, unbetont, nichtsyllabisch	m, n, r, l
Nasale und Liquiden, unbetont, syllabisch	$\underset{\circ}{m}, \underset{\circ}{n}, \underset{\circ}{r}, \underset{\circ}{l}$
Hohe Vokale, unbetont, nichtsyllabisch	y, i
Hohe Vokale, unbetont, syllabisch	u, i
Hohe Vokale, betont, syllabisch	$\acute{u}, \acute{i}$
Tiefe Vokale, unbetont	$o, e, a, \bar{o}, \bar{e}, \bar{a}$
Tiefe Vokale, betont	$\acute{o}, \acute{e}, \acute{a}, \acute{\bar{o}}, \acute{\bar{e}}, \acute{\bar{a}}$
Schwa secundum	$\circ$
leere Zeichenkette	$\square$

Ferner wird in Anlehnung an Rix et al. (1998:33) und Cowgill, W. & Mayrhofer, M. (1986:178f.) zu Lautgruppenzeichen zusammengefaßt.

$$R := \{m, n, r, l\}$$

$$\underset{\circ}{R} := \{\underset{\circ}{m}, \underset{\circ}{n}, \underset{\circ}{r}, \underset{\circ}{l}\}$$

$$H := \{h_1, h_2, h_3\}$$

$$\underset{\circ}{H} := \{\underset{\circ}{h}_1, \underset{\circ}{h}_2, \underset{\circ}{h}_3\}$$

$$K := \{\hat{k}, k, k^u, \hat{k}^h, k^h, k^{uh}, \hat{g}, g, g^u, \hat{g}^h, g^h, g^{uh}\}$$

$$T := \{t, t^h, d, d^h\}$$

Bei der Verwendung dieses Alphabets zur Schreibung nur von Teilmorphemkomplexen von Wortformen verlieren die durch die Symbolik nahegelegten Angaben zur Syllabizität im allgemeinen ihre Bedeutung, da Syllabizität wegen der ihr zugrundeliegenden Regeln ohne Einschränkungen höchstens für solche Morphemkomplexe sicher feststellbar ist, die vollständige Wortformen ergeben.

Wie auch aus dem obigen Zeicheninventar ersichtlich ist, wird der dreierlaryngalistische Ansatz vertreten. Hieraus kann mit geringem Aufwand eine nichtlaryngalistische Darstellung berechnet werden.

Grundsätzlich folgt der Programmieransatz dem im Vorwort von Rix et al. (1998:1-30) bzw. Meier-Brügger (2002:166ff.) und Rix (1992:190ff.) vorgestellten Modell zur Morphemauftrennung indogermanischer Verbalformen. Die maximale Anzahl der für die Bildung einer Verbalform anzusetzenden Morpheme beträgt derzeit sieben; es sind dies

Präfix (*meh*₁ für Prohibitiv bzw. Augment *h*₁*e* für Vergangenheitsmarkierung)

Reduplikationssilbe

Wurzel

Infix

Primärsuffix

Sekundärsuffix

Endung

In der Benutzeroberfläche sind diesen Morphemen unterschiedliche Farben zugeordnet, die in der beigegefügtten Abbildung nur als Graustufen erscheinen; in diesen Graustufen gleichartig hinterlegte Felder beinhalten unterschiedliche Informationen zum gleichen Morphem.

Der Vokalismus bzw. das Ablautverhalten und die Betonung der einzelnen Morpheme wird ausgerichtet an den Beschreibungen der 32 Primärstammbildungstypen aus Kapitel 3 des Vorwortes von Rix et al. (1998), die sämtlich ausgewertet wurden. Die realisierten Kombinationen aus Vokalismus, Ablaut- und Betonungsschemata korrespondieren nicht ganz mit den gegebenen Primärstammbildungstypen, denn die LIV-Typen 1a und 2a sowie 1f und 2b lassen sich solcherart, ohne weitere noch zu spezifizierende Attribute, nicht voneinander trennen.

Da vereinzelt Verbalwurzeln mit Wurzelvokal *a* und Nominalwurzeln mit Wurzelvokal *o* oder *a* in der Vollstufe vorkommen, sind *o*, *e*, *a* als ablautfähige Vokale einzustufen. Wegen des Morphems *eje* zur Bildung der LIV-

Typen 1s und 4a wird vereinbart, daß derjenige unter den Vokalen *o*, *e*, *a* eines Morphems abgelautet wird, der am weitesten rechts steht. Für die Belange der verbalen Formenbildung wurde bisher kein Morphem gefunden, das dieser Regel nicht genügen würde; alternativ wäre aber an eine weitergehende Morphemauftrennung zu denken, die dem Grundsatz folgt, daß jedes Morphem höchstens eine potentiell ablautfähige Stelle aufweist.

Der Wechsel beim Themavokal *e* wird als abhängig von der Paradigmastelle angesehen, da sich für einen Einfluß der lautlichen Umgebung keine schlüssige Regel aufstellen läßt. Wegen der Lückenhaftigkeit der Theorie werden im vorliegenden Programm keine Ablauterscheinungen bei Endungen modelliert; evtl. muß zu diesem Zweck hier gleichfalls eine Aufteilung in weitere Morpheme erwogen werden.

Die Betonung einer Wortform wird nach Abarbeitung des Ablauts und der Syllabierungsregeln auf den ersten Vokal eines Morphems gesetzt, wobei vorerst nur *o*, *e*, *a* und deren Ablautergebnisse sowie syllabische *u*, *i* berücksichtigt werden. Die syllabischen Nasale und Liquiden *m*, *n*, *ɲ*, *l* und die syllabischen Allophone der Laryngale *h*₁, *h*₂, *h*₃ werden als nicht betonungsfähig erachtet; diese Meinung muß vielleicht wenigstens hinsichtlich der Nasale und Liquiden revidiert werden.

Die etwa 1200 Wurzeln wurden aus LIV übernommen, wobei durch Klammern angezeigte Alternativen aufgelöst wurden; die Fußnoten blieben bislang unausgewertet, ebenso die 2001 erschienene 2. Auflage. Der inkonsistente Standpunkt des LIV hinsichtlich des Auftretens von *Tenues aspiratae* im Indogermanischen verrät eine zeitlich inhomogene Schichtung des gebotenen Materials, dem dadurch Rechnung getragen wurde, daß zu allen Wurzeln mit *Tenues aspiratae*, bei denen der Lautwandel nach Siebs rückgängig gemacht werden kann, auch Alternativen mit *Mediae aspiratae* in die Verbalformenwerkbank aufgenommen wurden.

Die Endungssätze stammen aus Beekes (1995), Meier-Brügger (2002), Meiser (1998), Rix (1992) und Tichy (2000).

Unter Bezugnahme auf Rix (1992: Par. 216) und die in Kapitel 3 des Vorwortes von Rix et al. (1998) aufgeführten Formeln zur Reduplikation werden als Reduplikationstypen bereitgestellt

Abk.	Reduplikationssilbe
NUL	□
SAN	wird gebildet aus dem ersten Anglittradikal der Wurzel verkettet mit dem zweiten Anglittradikal der Wurzel, falls ein zweites Anglittradikal vorhanden ist und der erste Anglittradikal gleich <i>s</i> ist, verkettet mit dem separat anzugebenden Reduplikationsvokal.
ANG	wird gebildet aus dem Anglitt der Wurzel verkettet mit dem separat anzugebenden Reduplikationsvokal.
WRZ	wird gebildet aus dem Anglitt der Wurzel verkettet mit dem separat anzugebenden Reduplikationsvokal verkettet mit dem Abglitt der Wurzel.
ANB	wird gebildet aus dem Anglitt der Wurzel verkettet mit dem separat anzugebenden Reduplikationsvokal verkettet mit dem ersten Abglittradikal der Wurzel.

Weitere Reduplikationsprinzipien sind denkbar; vgl. Beekes (1995:227). Der Reduplikationstyp ANG wurde lediglich zu experimentellen Zwecken geschaffen, alle anderen findet man mittelbar oder unmittelbar auch bei den LIV-Typen. Der Reduplikationstyp SAN wird gebraucht bei den LIV-Typen 1g, 1h, 1i, 2c, 3a, 5b, ANB bei LIV-Typ 6a, und WRZ bezieht sich auf die Intensiv-Reduplikation nach (Rix 1992: Par. 216).

Die in LIV angegebenen Formeln zur Reduplikation sind jedoch im Hinblick auf zweistellige Anglittes formal unsauber und wurden daher entsprechend modifiziert. Würde man die zu (6a) gehörende Formel streng umsetzen, entstünde zur Wurzel $g^u\hat{g}^her$ der Intensivstamm $g^u\hat{e}\hat{g}^h-g^u\hat{g}^hor/g^u\hat{g}^hr$ - und nicht $g^u\hat{g}^hér-g^u\hat{g}^hor/g^u\hat{g}^hr$ -, wie in Rix et al. (1998:190) mit Fragezeichen angemerkt. Bei den anderen Typen mit nichtleerer Reduplikationssilbe bliebe Rix (1992: Par. 216) unbeachtet; zur Wurzel *steh*₂ entstünde der Perfektstamm *se-stóh*₂/*sth*₂- und nicht *ste-stóh*₂/*sth*₂- (Rix et al. 1998:536).

Bei der Betrachtung grammatischer Kategorien wird deutlich, daß für eine kohärente, ohne Ausnahmeregelungen auskommende und unter einheitlichen Gesichtspunkten durchgeführte Programmierung für die gesamte verbale Formenvielfalt eine Modifikation der Sichtweise auf die Bestimmungsgrößen der Paradigmen erforderlich ist.

Zwischen den grammatischen Kategorien Tempus, Modus, Diathese, Aspekt und Aktionsart ist in der befragten sprachwissenschaftlichen Literatur keine einheitliche scharfe Trennung vorgenommen, vielleicht auch, weil dies nicht ohne weiteres möglich ist. Ebensowenig ergibt sich ein klares Bild hin-

sichtlich der für die jeweiligen Zeitstufen anzusetzenden Kombinationen von Kategorien. Es bleibt also unklar, ob oder wann es beispielsweise Konjunktiv Perfekt bzw. Konjunktiv Stativ gibt oder Medialformen beim LIV-Typ 3a.

Mit *formaler Modus* wird im weiteren Tempus-Modus verkürzt und von einigen indogermanischen Vorstellungen geringfügig abweichend wiedergegeben, bei Diathese werden, angeregt von der vorurindogermanischen Sichtweise mancher Autoren, z.B. Tichy (2000:91), die Attributsausprägungen Perfekt und Stativ mit eingeschlossen.

Die der Datenbank zugrundegelegten Auffassungen der formalen Modi unterscheiden sich von den etablierten indogermanistischen Lehrmeinungen auch insofern, als jedem solchen Modus eine Kombination aus Zeichenkette im Augment, ggf. ablautender Zeichenkette im Sekundärsuffix und Typ des Endungssatzes zugeordnet ist. Damit sind die unter formalen Modi zusammengefaßten Attributsausprägungen Indikativ, Imperfekt, Optativ und Konjunktiv eindeutig charakterisiert; Imperativ und Injunktiv können zusammenfallen.

formaler Modus	Zeichenkette im Augment	Zeichenkette im Sekundärsuffix	Typ des Endungssatzes
Indikativ	□	□	PE
Injunktiv	□	□	SE
Imperfekt	<i>he</i>	□	SE
Optativ	□	<i>ieh₁</i>	SE
Konjunktiv	□	<i>he</i>	PE oder SE
Imperativ	□	□	SE oder SE modifiziert

In der Tabelle steht □ für die leere Zeichenkette, PE für primären Endungssatz, SE für sekundären Endungssatz; die für Augment und Sekundärsuffix gewählten Zeichenketten stammen aus Tichy (2000).

Die formale Auffassung wirkt sich lediglich störend aus bei den LIV-Typen 2a, 2b, 2c und 3a, also bei Aorist und Perfekt, hat aber neben der strengeren Systematik den Vorteil, daß man sich auch die Möglichkeit zur Modellierung bislang unbeachtet gebliebener Kombinationen von Suffixen und Endungssätzen offenhält. In der nächsten Tabelle werden die beiden Auffassungen einander gegenübergestellt.

LIV-Typ	formaler Modus	indogermanistische Sichtweise
2a, 2b, 2c	Indikativ	existiert nicht
2a, 2b, 2c	Imperfekt	Indikativ
3a	Imperfekt	Plusquamperfekt

Berücksichtigte Regeln zur Lautung sind die Syllabierungsregeln einschließlich der Regeln von Sievers und Lindeman. Hierbei dient als Leitlinie grob die untenstehende, aus Gippert (2001: Par. 4.2.1.) entnommene Hierarchie. Weitere Lautwandelprozesse bleiben vorerst außer Betracht.

Rank	Sound category	Conditions of syllabicity
I.	Non-high vowels	Always (wherever they appear)
II.	High vowels, nasals, liquids	When not adjacent to sound of category I and when not followed by syllabic sound of category II ("Schindler's rule")
		when, as second element of syllable, followed by sound of category I in word-final syllable ("Sievers" and "Lindeman's" laws)
III.	Laryngeals	When not adjacent to sound of category I or syllabic sound of category II
IV.	"Shewa secundum"	Within remaining clusters of non-syllabic consonants (facultatively)

Zusätzlich einbezogen wurden nach Cowgill & Mayrhofer (1986: Par. 7.3.2.f.) im Wurzel-Infix-Komplex Bedingungen zur Syllabierung der Gruppen μR , mR und des Infixes n der Nasalpräsentien (LIV-Typ 1k); ansonsten wird Schindlers Regel (Schindler 1977) wie vorgeschrieben von rechts nach links abgearbeitet.

Sowohl die Regel von Sievers als auch die von Lindeman wurden ausgedehnt auf Nasale und Liquiden; es sind also die Varianten Sievers-Edgerton und Lindeman-Edgerton implementiert.

Schwa secundum entsteht in der derzeitigen Programmfassung lediglich zwischen zwei Obstruenten im aus Wurzel und Infix gebildeten Morphemkomplex.

3.2. Die Benutzeroberfläche

Ausw. 1		LIV-Typ	Diathese	Modus	Bedung nach	Status 1	LIV-Typ	(1a)	Bedung nach	Tichy
(1a) ▾		Vorwahl: 1a	Alt. ▾	* ▾	Tichy ▾	Mod u. Dia.		Ind. Alt.	Wurzelsoder	Handlungsbe
Ausw. 2		Augmento.Faß	Reduplikationsbe	Wurzel	Inf	Primärsuffix		Sekundärsuffix	Bedung	
Morph.-bel		▾	▾	hɛj	▾	▾		▾	▾	▾
Abl./Vok.		▾	▾	▾	▾	▾		▾	▾	▾
Bedung		▾	▾	▾	▾	▾		▾	▾	▾

Status 2	Morphembelegung	Ablaut/Vokalismus										Bedung						
1psng	▾	NUL	hɛj	▾	▾	mi	a	e	a	a	a	e	▾	'	▾	▾	▾	▾
2psng	▾	NUL	hɛj	▾	▾	si	a	e	a	a	a	e	▾	'	▾	▾	▾	▾
3psng	▾	NUL	hɛj	▾	▾	ti	a	e	a	a	a	e	▾	'	▾	▾	▾	▾
1psdu	▾	NUL	hɛj	▾	▾	ues	a	a	a	a	a	e	▾	'	▾	▾	▾	▾
2psdu	▾	NUL	hɛj	▾	▾	thɛɔs	a	a	a	a	a	e	▾	'	▾	▾	▾	▾
3psdu	▾	NUL	hɛj	▾	▾	tes	a	a	a	a	a	e	▾	'	▾	▾	▾	▾
1pspl	▾	NUL	hɛj	▾	▾	mes	a	a	a	a	a	e	▾	'	▾	▾	▾	▾
2pspl	▾	NUL	hɛj	▾	▾	tes	a	a	a	a	a	e	▾	'	▾	▾	▾	▾
3pspl	▾	NUL	hɛj	▾	▾	enti	a	a	a	a	a	e	▾	'	▾	▾	▾	▾

Formen	Singular	Dual	Plural
1. Person	hɛjmi	hɛjɛs	hɛjmɛs
2. Person	hɛjsi	hɛjthɛɔs	hɛjɛs
3. Person	hɛjni	hɛjɛs	hɛjenti

Paradigma zur Wurzel *hɛj* 'gehen' (Rix et al. 1998:207f.), LIV-Typ 1a (amphidynamisches Wurzelpresens); vgl. Paradigma in Szemerényi (1990:343)

Es wird nun kurz auf die Zweckbestimmung der Bereiche *Auswahl 1* (*Ausw. 1*), *Auswahl 2* (*Ausw. 2*), *Status 1*, *Status 2*, *Formen* eingegangen. Die Felder des Bereiches *Formen* übernehmen die Datenausgabe; für sie wird, neben einigen weiteren Hilfszeichen zur Bewertung der angezeigten Ergebnisse, das eingangs in diesem Kapitel vorgestellte Alphabet in vollem Umfang verwendet. Die Felder der anderen Bereiche sind der Dateneingabe und Datenmodifikation sowie der Kontrolle über die gewählten Einstellungen nach grammatischen Klassifikationsschemata und unterschiedlichen Lehrmeinungen vorbehalten. Als Eingangsgrößen für die Morpheme werden, für jedes Morphem einzeln, separat die Gestalt der Zeichenkette seiner Vollstufe, die Ablautstufe und die Betonung angegeben.

Für die Angabe der Gestalt der Zeichenkette werden somit keine Zeichen für betonte Vokale benötigt; für die Interpretation hinsichtlich Syllabizität gilt das im Einleitungsteil zu diesem Kapitel Gesagte.

Zur Bezeichnung der Ablautstufe dienen

leerer Vokalismus	□
Schwundstufe	ə
Vollstufe	e
abgetönte Vollstufe	o
Dehnstufe	ē
abgetönte Dehnstufe	ō

wobei diese Zeichen als symbolische Größen zu verstehen sind, mit denen per se noch kein Lautwert explizit gemacht werden soll, und zur Bezeichnung der Betonung

unbetont	°
betont	ˈ

Bereich *Auswahl 1*

In den vier rechts neben dem Schriftzug *Vorwahl* stehenden Feldern kann man Parameter für das Auswahlfeld ganz links konkret voreinstellen. In das mit LIV-Typ beschriftete Textfeld werden geeignete, von a bis v, von 1 bis 8 reichende Ziffern- und Buchstabenkombinationen eingetragen (vgl. Rix et al. 1998), unter Diathese kann man sich entscheiden zwischen Aktiv, Medium,

Perfekt, Stativ, unter Modus zwischen Imperativ, Indikativ, Injunktiv, Konjunktiv, Optativ, und unter Endung werden die Endungssätze nach Beekes, Meier-Brügger, Meiser, Rix, Tichy angeboten. Daneben ist es für jedes derartige Feld möglich, die Voreinstellung durch den Kleene-Stern * unspezifiziert zu lassen.

Im Auswahlfeld ganz links erscheinen nun, abhängig von der getroffenen Vorwahl, die gemäß LIV und den aus den befragten Lehrbüchern gewonnenen theoretischen Informationen maximal denkbaren Tempus-Modus-Aspekt-Aktionsart-Kombinationen.

Hierbei kann man für den Konjunktiv zwischen Sekundär- oder Primärendungssatz wählen, und bei den Imperativen, ggf. mit alternativen Belegungen für die Endungen gewisser Paradigmastellen, zwischen Imperativ I oder Imperativ II.

Nachdem man seine Wahl endgültig getroffen hat, werden alle nicht die Zeichenkettengestalt der Wurzel betreffenden Felder der Bereiche *Status 1* und *Status 2* gefüllt.

Bereich *Status 1*

Die Felder beinhalten weitgehend Informationen für den Benutzer und zeigen die beim Auswahlvorgang in *Auswahl 1* eingestellten Parameter zu grammatischen Klassifikationsschemata und unterschiedlichen Lehrmeinungen an, ferner die noch in *Auswahl 2* festzulegende Gestalt der Wurzel, sofern nicht schon unmittelbar hier in das betreffende Feld eine Wurzel eingegeben wird.

Bereich *Auswahl 2*

Für die sieben in die Bildung finiter Wortformen einzubeziehenden Morpheme wird, im Sinne der oben besprochenen Dreiteilung der Eingangsgrößen, in drei Zeilen von Auswahlfeldern Zugriff auf alle gängigen Vollstufen von Morphemen, auf Ablaut- und Betonungsschemata gewährt. Die obere Zeile ist der Zeichenkettendarstellung der Morpheme vorbehalten, wobei im dritten Feld von links die in LIV verzeichneten indogermanischen Wurzeln abgerufen werden können.

Mit einer hier getätigten Auswahl können vorher getroffene Einstellungen modifiziert und ergänzt werden; diese Änderungen werden in den Feldern des Bereiches *Status 2* sichtbar.

Bereich *Status 2*

Die Angaben in den Feldern dieses Bereiches werden unmittelbar als Eingabeparameter für den Formenerzeugungsalgorithmus genommen. Falls der Benutzer es wünscht, können in jedem einzelnen Textfeld, also für jede Paradigma-stelle und für jede irgendein Morphem betreffende Information, im Rahmen der vereinbarten Notationsgepflogenheiten noch manuell Änderungen vorgenommen werden.

Bereich *Formen*

Dieser Bereich enthält die Ausgabefelder des Wortformenalgorithmus. Die Bedeutung der Felder ergibt sich aus den Beschriftungen am oberen und linken seitlichen Rand des Bereichs.

3.3. *Beschreibung des Leistungsumfangs der Verbalformenwerkbank*

Ohne Zuhilfenahme einer Klassifikation zur Formenbildung aus Verbalwurzeln, wie sie in beispielhafter Weise durch LIV (Rix et al. 1998) vorgelegt ist, scheint eine umfassende Überprüfung der durch die indogermanischen Verbalformenwerkbank gelieferten Resultate kaum möglich. Eine naheliegende Vorgehensweise ist es, zu einer beliebigen Verbalwurzel das Äquivalent in LIV aufzusuchen, die entsprechenden, nur indirekt durch Stammansätze vermittelten LIV-Typen durch Nachschlagen im Einleitungsteil von Rix et al. (1998:14-25) zu bestimmen, und dann Paradigmen generieren zu lassen.

Für die Wurzel *b^hérǵ^h* ‘hoch werden, sich erheben’ (Rix et al. 1998:63) gelangt man zu der Übersicht

Stammansatz	Aspekt-Aktionsart	LIV-Typ
<i>b^hérǵ^h/b^hrǵ^h-</i>	Aorist	2a
<i>b^hrǵ^h-jé-</i>	Präsens	1q
<i>b^horǵ^h-éje-</i>	Kausativ	4a
<i>b^he-b^hórǵ^h/b^hrǵ^h-</i>	Perfekt	3a

und für *g^uhen* ‘schlagen’ (Rix et al. 1998:194)

Stammansatz	Aspekt-Aktionsart	LIV-Typ
$g^{uhén-}/g^{uhn-}$	Präsens	1a
$g^{uhj-}g^{uhn-é-}$	Präsens	1i
$g^{uhn-ské-}$	Präsens	1p
$g^{uhé-}g^{uhn-e-}$	Aorist	2c
$g^{uhon-éje-}$	Iterativ	4a

Es werden nun mit den Grundeinstellungen der indogermanischen Verbalformenwerkbank unter Zuhilfenahme der Angaben in LIV einige Paradigmen erzeugt, denen die in der Literatur ausgeführten Paradigmen gegenübergestellt werden. Leider sind dort nur spärlich Exemplare anzutreffen; insgesamt wurden etwa 25 Vergleichsparadigmen aus Beekes (1995), Szemerényi (1990) und Tichy (2000) entnommen, wo erwartungsgemäß paarweise voneinander verschiedene theoretische Auffassungen vertreten werden, die häufiger auch nicht so recht mit den in Rix et al. (1998:14-25) beschriebenen Stammbildungstypen harmonieren wollen. Zur Einschätzung der Qualität der Verbalformenwerkbank gehen in die nachfolgende Auswertung die aus allen diesen Paradigmenvergleichen gewonnenen Erkenntnisse ein; drei dieser Paradigmenvergleiche werden anschließend im Detail besprochen.

Die Differenzen zwischen den automatisch erzeugten und den in Druckwerken gegebenen Paradigmen werden möglichst penibel bewertet und durch Hinterlegung gekennzeichnet; die Kennzeichnung folgt starr dem Prinzip, daß bei einander als entsprechend gedachten Wortformen jeder Unterschied markiert wird, der sich in auch nur einem Zeichen niederschlägt. Selbstverständlich werden hierbei die von den Autoren bevorzugten Normalisierungen berücksichtigt; Segmentierungsstriche werden nicht in die Betrachtung einbezogen, aber durch Bindestriche eingeleitete Bequemlichkeitsschreibweisen sowie durch Klammerung angezeigte Optionalität oder gar Fragezeichen werden als Interpretationsspielräume behandelt, die dann dem Gutdünken ausgeliefert sind, wenn sich ein vom Autor intendierter Rahmen nicht mehr klar erkennen läßt.

Jedem Paradigmenvergleich geht eine Diskussion voran über die in der indogermanischen Verbalformenwerkbank getroffenen Einstellungen und den gedanklich mutmaßlich zugrundezulegenden Standpunkt des Autors; es folgt dem Paradigmenvergleich eine dezidierte Analyse der aufzufindenden Unterschiede.

3.3.1. Tichy (2000:103f.)

Beispielparadigmen zu *g^uhen* ‘schlagen’ (Rix et al. 1998:194), LIV-Typ 1a (amphidynamisches Wurzelpräsens)

Imperfekt Aktiv

Einstellung Verbalformenwerkbank

a. Status 1

LIV-Typ	Mod. u. Dia.	Endung nach	Wurzel
1a	Ipf. Akt.	Tichy	<i>g^uhen</i>

b. Status 2

keine Modifikationen

Standpunkt Tichy: Vermutlich keine Abweichung zur Verbalformenwerkbank

Paradigmenvergleich

Paradigmastelle	Verbalformenwerkbank	Tichy
1psng	<i>h₁ég^uhenm</i>	<i>é-g^{wh}en-m</i>
2psng	<i>h₁ég^uhens</i>	<i>é-g^{wh}en-s</i>
3psng	<i>h₁ég^uhent</i>	<i>é-g^{wh}en-t</i>
1psdua	<i>h₁ég^uh_ηue</i>	<i>é-g^{wh}η-ue</i>
2psdua	<i>h₁ég^uh_ηteh₂</i>	<i>é-g^{wh}η-tah₂</i>
3psdua	<i>h₁ég^uh_ηteh₂m</i>	<i>é-g^{wh}η-tah₂m</i>
1psplr	<i>h₁ég^uh_ηmo</i>	<i>é-g^{wh}η-me</i>
2psplr	<i>h₁ég^uh_ηte</i>	<i>é-g^{wh}η-te</i>
3psplr	<i>h₁ég^uh_ηnt</i>	<i>é-g^{wh}η-ent</i>

Analyse der aufgefundenen Unterschiede

- Alle Paradigmastellen: *h₁é* vs. *é*; unterschiedliche Auffassungen über die Gestalt des Augments
- 2psdua, 3psdua: *e* vs. *a*; Laryngalumfärbung nicht vollzogen vs. Laryngalumfärbung vollzogen
- 1psplr, 3psplr: *o* vs. *e*, *η* vs. *en*; unterschiedliche Ablautstufen in der Endung

3.3.2. Beekes (1995:257)

Beispielparadigmen zu k^uej ‘sammeln’ (Rix et al. 1998:338), LIV-Typ 11 (Präsens mit Suffix $néu/nu$)

Imperfekt Aktiv

Einstellung Verbalformenwerkbank

a. Status 1

LIV-Typ	Mod. u. Dia.	Endung nach	Wurzel
11	Ip. Akt.	Beekes	k^uej

b. Status 2

keine Modifikationen

Standpunkt Beekes: Vermutlich keine Abweichung zur Verbalformenwerkbank

Paradigmenvergleich

Paradigmastelle	Verbalformenwerkbank	Beekes
1psng	$h_1ék^uej\bar{m}$	$h_1ék-k^uej-neu-m$
2psng	$h_1ék^uej\bar{s}$	-s
3psng	$h_1ék^uej\bar{t}$	-t
1psdua	$h_1ék^uej\bar{nu}e$	
2psdua	$h_1ék^uej\bar{nu}tom$	
3psdua	$h_1ék^uej\bar{nu}teh_2\bar{m}$	
1psplr	$h_1ék^uej\bar{nu}me$	-nu-mé
2psplr	$h_1ék^uej\bar{nu}te$	-té
3psplr	$h_1ék^uej\bar{nu}tent$	-(e)nt

Analyse der aufgefundenen Unterschiede

- 1) 1psng: $\bar{m}$ vs. um ; unterschiedliche Auffassungen zur Syllabierung
- d) 1psplr, 2psplr: e vs. $é$; widersprüchliche Auffassungen von Beekes hinsichtlich der Betonung; siehe Beekes (1995:226) ‘... the augment took the stress’

3.3.3. Szemerényi (1990:342f.)

Beispielparadigmen zu *h₁es* ‘dasein, sein’ (Rix 1998:214f.), LIV-Typ 1a (amphidynamisches Wurzelpräsens)

Optativ Aktiv

Einstellung Verbalformenwerkbank

a. Status 1

LIV-Typ	Mod. u. Dia.	Endung nach	Wurzel
1a	Opt. Akt.	Tichy	<i>h₁es</i>

b. Status 2

keine Modifikationen

Standpunkt Szemerényi: Wurzel *es*, ansonsten nicht ersichtlich

Paradigmenvergleich

Paradigmastelle	Verbalformenwerkbank	Szemerényi
1psng	<i>h₁sjéh₁m</i>	<i>syēm</i>
2psng	<i>h₁sjéh₁s</i>	<i>syēs</i>
3psng	<i>h₁sjéh₁t</i>	<i>syēt</i>
1psdua	<i>h₁sih₁ué</i>	
2psdua	<i>h₁sih₁téh₂</i>	
3psdua	<i>h₁sih₁téh₂m</i>	
1psplr	<i>h₁sih₁imé</i>	<i>sīme</i>
2psplr	<i>h₁sih₁ité</i>	<i>sīte</i>
3psplr	<i>h₁sih₁ént</i>	<i>sīyent</i>

Analyse der aufgefundenen Unterschiede

- 1) 1psng, 2psng, 3psng, 1psplr, 2psplr, 3psplr: *h₁s* vs. *s*; unterschiedliche Auffassungen zur Laryngaltheorie
- e) 1psng, 2psng, 3psng: *éh₁* vs. *ē*; unterschiedliche Auffassungen zur Laryngaltheorie, keine explizite Notierung der Betonung bei Szemerényi
- f) 1psplr, 2psplr, 3psplr: *ih₁* vs. *ī*, *ih₁é* vs. *īye*; unterschiedliche Auffassungen zur Laryngaltheorie

- g) 1psplr, 2psplr, 3psplr: *é* vs. *e*; fehlende Notierung der Betonung bei Szemerényi
- h) 1psng: *m̄* vs. *m*; unterschiedliche Durchführungen der Syllabierung

3.4. Auswertung der Abweichungen

In die Auswertung wurden sämtliche bisher angestellten Paradigmenvergleiche einbezogen; nicht alle der unten protokollierten systematisch auftretenden Unterschiede finden sich in den in dieser Arbeit vorgestellten Paradigmenvergleichen.

3.4.1. Fazit zu Tichy

Es sind drei Typen systematisch auftretender Unterschiede zu erkennen.

- 1) Laryngalumfärbung nicht vollzogen vs. Laryngalumfärbung vollzogen
- i) Unterschiedliche Auffassungen von der Gestalt des Augments; siehe hierzu Tichy (2000: Par. 13.3.1, Par. 16.2.2)
- j) Unterschiedliche Auffassungen von den zu verwendenden Endungsallomorphen

Die Ausführungen in Tichy (2000: Par. 11.3.1.), bestehend aus einer tabellarischen Auflistung der möglichen Endungsallomorphe und begleitenden Kommentaren, ergibt meines Erachtens kein vollständiges, in sich schlüssiges Bild. Hilfreich zur Findung eines Gesamtkonzeptes ist (Rix 1992: Par. 265), *‘Ablautvarianten sind durch vergleichende Rekonstruktion nur für die Endungen der 3. Pl. wie -enti/nti und für -mes/mos der 1. Pl. zu fassen. Sie entstammen verschiedenen Akzent-Ablaut-Paradigmen [...]: e-Stufe setzt Betonung der Plur.-Endungen voraus, o- oder Schwundstufe durchgängige Betonung von Wurzel oder Primäraffix.’* Dieses Prinzip wurde den Endungen nach Tichy bei der Programmierung der indogermanischen Verbalformenwerkbank zugrundegelegt.

Insgesamt erweisen sich die Differenzen zwischen Beispielparadigmen aus Tichy und den von der Verbalformenwerkbank generierten Paradigmen als so wenig gravierend, daß eine völlige gegenseitige Übereinstimmung leicht erreichbar scheint.

3.4.2. *Fazit zu Beekes*

Es sind zwei Typen systematisch auftretender Unterschiede zu erkennen.

- 1) Unterschiedliche Auffassungen zur Syllabierung
- k) Unterschiedliche Auffassungen bei Ablaut- und Betonungsschemata

Die Gepflogenheiten hinsichtlich der Markierung von Syllabizität sind in Beekes (1995:125) dargelegt; demgemäß hält Beekes eine Unterscheidung syllabischer und nichtsyllabischer Resonanten durch die Symbolik für nicht erforderlich. Eine völlige Übereinstimmung mit den erzeugten Paradigmen ist möglicherweise erreichbar, wenn es gelingt, die in Beekes (1995) verstreut auffindbaren Auffassungen zu Ablaut und Betonung zu komplettieren und als Grundeinstellungen in den Formenerzeugungsalgorithmus zu implementieren; im jetzigen Stadium lehnen sich die der indogermanischen Verbalformenwerkbank zugrundeliegenden Ablaut- und Betonungsschemata noch weitestgehend möglich an LIV an. Es wäre hier also die eingangs angemahnte klare Abgrenzung konkurrierender Denkmodelle voneinander vorzunehmen und der Unterschied zwischen den in Beekes und in LIV vertretenen Standpunkten explizit zu machen.

Jenseits der vorgegebenen Grundeinstellungen kann man sich schon mit der derzeitigen Version der indogermanischen Verbalformenwerkbank durch gezielte Auswahl in den Feldern des Bereiches *Auswahl 2* oder durch manuelle Änderungen in den Feldern des Bereiches *Status 2* den Wortformen nach Beekes weiter annähern.

Insgesamt erweisen sich die Differenzen zwischen den bisher betrachteten Beispielparadigmen aus Beekes und den von der Verbalformenwerkbank generierten Paradigmen aber als wenig gravierend.

3.4.3. *Fazit zu Szemerényi*

Es sind mindestens 5 Typen systematisch auftretender Unterschiede zu erkennen.

- 1) Unterschiedliche Auffassungen zur Laryngaltheorie
- l) Unterschiedliche Ablautstufen in den Endungen oder unterschiedliche Morpheme für die Endungen
- m) Unterschiedliche Ablautstufen in Wurzeln oder Suffixen
- n) Unterschiedliche Auffassungen zur Syllabierung

o) Unterschiedliche Betonung

In Szemerényi wird dargelegt, daß einem nichtlaryngalistischen Ansatz der Vorzug gegeben wird. Die Endungen nach Szemerényi wurden nicht in die indogermanischen Verbalformenwerkbank eingearbeitet, da die Angaben zu lückenhaft waren; stattdessen wurde bei der Generierung der Beispiele aus Szemerényi durch die Verbalformenwerkbank auf die Endungssätze nach Tichy zurückgegriffen. Trotzdem ist die Abweichungsquote erstaunlich gering, wenn man zusätzlich noch die schon angemahnte Inkonsistenz bei den Ablautstufen der Endungssätze nach Tichy in Rechnung stellt. Die Bildung von Verbalstämmen nach Szemerényi und nach LIV und die zugehörigen Betonungs- und Ablautmuster unterscheiden sich erheblich; wie bei Beekes müßte auch hier eine genauere Analyse der Bildung von Verbalstämmen nach Szemerényi erfolgen sowie ggf. eine klare Abgrenzung der beiden konkurrierenden Denkmodelle voneinander betrieben werden.

Die Differenzen hinsichtlich der Syllabierung und der Betonung gehen meines Erachtens, wie bereits am Ende von Kapitel 2 erwähnt, eindeutig zu Lasten von Szemerényi, der nach nicht nachvollziehbaren Kriterien an unterschiedlichen Stellen zu unterschiedlichen Ausdrucksmitteln greift.

Für eine weitere Annäherung der von der derzeitigen Version der indogermanischen Verbalformenwerkbank jenseits der vorgegebenen Grundeinstellungen generierbaren Wortformen an die Wortformen nach Szemerényi gilt analog das oben zu Beekes Gesagte.

Insgesamt erweisen sich die Differenzen zwischen den bisher betrachteten Beispielparadigmen aus Szemerényi und den von der Verbalformenwerkbank generierten Paradigmen als nicht unüberbrückbar.

Die Disambiguierung der Angaben in den befragten Lehrbüchern war bisweilen schwierig, manchmal unmöglich, gleichfalls selbstredend die Abstimmung unterschiedlicher Theoriebildungen aufeinander oder die Normierung auf eine gemeinsame Zeitstufe. Daher entstehen bei der Formenerzeugung mit ziemlicher Sicherheit sowohl zeitlich als auch hinsichtlich der Theoriebildungen hybride Lösungen und inhomogene Paradigmen, wenn man etwa an die Endungssätze nach Beekes ohne die detaillierte Kenntnis der von ihm bevorzugten Akzent- und Ablautschemata oder die unterschiedlichen Sichtweisen zu Schindler, Sievers, Lindeman und Edgerton denkt; gerade für den Komplex aus Syllabierungsregeln und Sievers-Lindeman-Edgerton gibt es gewiß noch

einiges an Korrekturen vorzunehmen, wobei weitere Ansätze, etwa nach Keydana (2004), einzubeziehen wären.

Hinsichtlich des Vergleiches der generierten Formen mit denen aus der Literatur (Beekes 1995; Szemerényi 1990; Tichy 2000) ist die Stichprobe mit etwa 25 Exemplaren viel zu klein, um hieraus endgültige Aussagen treffen zu können. Die aufgefundenen Differenzen sind immerhin klassifizierbar und haben nachvollziehbare Ursachen; meist sind die Abweichungen geringfügig und aus den unterschiedlichen Theorien sowie durch Eigenheiten der Autoren und durch von ihnen zu verantwortende Inkonsistenzen und Widersprüchlichkeiten erklärbar. Nach dem derzeitigen Stand der Dinge erscheint es realistisch, eine in sich schlüssige und mit den bestehenden Theorien korrelierende indogermanische Verbalformenwerkbank herstellen zu können, sobald diese Theorien widerspruchsfrei formuliert und voneinander sauber abgegrenzt sind.

Weiterhin sind die erzeugten Wortformen rein virtuelle Gebilde, die noch nicht hinreichend an die sprachliche Realität angeschlossen sind. Eine Aufnahme der Indizes zur rekonstruierten Stammbildung (Rix et al. 1998:649-661) und zu nachfolgesprachlichen Wortformen (Rix et al. 1998:663-754) in die Verbalformenwerkbank würde dies ändern und eine der möglichen Anwendungen der in Kapitel 1 vorgestellten Konzeption gestatten, wenn man, ausgehend von den durch die Verbalformenwerkbank generierten indogermanischen Wortformen, den Lautwandel bis zu den betreffenden Nachfolgesprachen hin modellieren würde.

Eine der Verbalformenerzeugung vergleichbare indogermanische Nominalformenerzeugung ist ein erreichbares Desideratum.

4. *Modellierung des Lautwandels vom Urgermanischen zum Altisländischen*

4.1. *Aufbau und Zielsetzung*

Ein Anfang 2003 fertiggestelltes Programm ermöglicht die Modellierung von Lautwandelprozessen unabhängig von Einzelsprachen und Beobachtungszeitraum. Zur Darstellung der Regeln der zu bearbeitenden Lautwandelerscheinungen und des zugehörigen Wortmaterials, die separat in Tabellen abgelegt sind und auf die das Programm zugreift, wurde eine Syntax entwickelt, die es erlaubt, eine endliche, jedoch beliebig hohe Anzahl von Merkmalsdimensio-

nen, wie beispielsweise Lautwert, suprasegmentale Merkmale oder Morphemgrenzen, zu berücksichtigen; dies schließt die Möglichkeit der binären Kodierung von Merkmalen mit ein. Zur Abarbeitung der Regeln können unterschiedliche Parameter eingestellt werden; neben der üblichen sequentiellen Abarbeitung ist es möglich, die Gültigkeitsdauer von Regeln über einen beliebigen Zeitraum hinweg einzustellen, Regeln iterativ anzuwenden, Prioritäten zu vergeben, oder die Durchlaufrichtung einer Regel zu vereinbaren. Ein Beispiel für eine Regel, bei der es auf die Durchlaufrichtung ankommt, ist die Regel von Schindler (Schindler 1977) zur Entscheidung über Syllabizität, die iterativ von rechts nach links abzuarbeiten ist.

Bereitgestellt werden sollen ferner Möglichkeiten, zwischen obligatorischen und fakultativen Regeln zu unterscheiden. Die Grenzen des Algorithmus liegen alleine in den Ressourcen der eingesetzten Rechnerumgebung.

Das Programm wurde bisher anhand von ca. 400 Wortformen aus Cathey (1967) erfolgreich getestet, wobei der letzte Testlauf im März 2003 erfolgte.

Es soll nun ein etwas tieferer Einblick gewährt werden in die der Modellierung des Lautwandels vom Urgermanischen zum Altisländischen zugrundegelegten Notationen, die sowohl innerhalb des Algorithmus als auch zur Dateneingabe und -ausgabe benutzt werden. Die Auswahl der Symbolik ist vorrangig auf linguistische Bedürfnisse ausgerichtet; trotz einiger Anleihen aus dem mathematischen Zeicheninventar wird an dieser Stelle aber darauf verzichtet, auf die zugrundeliegenden Strukturen unter Verwendung mathematischer Sprechweisen einzugehen. In der genauesten Darstellung wird jede lautliche Einheit, sei es die einer Wortform oder die einer Regel des Lautwandels, aufgefaßt als Element eines Merkmalsraumes, der aus den vier Merkmalsdimensionen Position im Wort, phonologische Jungiertheit, Betonung und Lautwert gebildet ist; mit $\oplus$ als Verknüpfungszeichen für die Merkmalsdimensionen wird also jede lautliche Einheit geschrieben als vierstelliger Ausdruck von Merkmalsausprägungen in der Form

$$\text{Position} \oplus \text{Jungiertheit} \oplus \text{Betonung} \oplus \text{Lautwert}$$

Für die Merkmalsausprägungen der Position im Wort dienen

Wortgrenze links, keine Wortgrenze rechts	←
Wortgrenze rechts, keine Wortgrenze links	→
Wortgrenze links und rechts	↔
Wortgrenze weder links noch rechts	↔

Für die Merkmalsausprägungen der phonologischen Jungiertheit, wie sie beispielsweise bei Polyphthongen vorkommt, aber auch Affrikaten zugesprochen werden könnte, werden verwendet

Jungiertheit nach links, keine Jungiertheit nach rechts	↔
Jungiertheit nach rechts, keine Jungiertheit nach links	↔
Jungiertheit nach links und nach rechts	↔
Jungiertheit weder nach links noch nach rechts	↔

Die Wortbetonung wird gemäß dem von Cathey präferierten, aus den drei Stufen *primary*, *secondary*, *tertiary stress* bestehenden Modell ausgezeichnet mit

Hauptbetonung	!!
Nebenbetonung	!
Unbetontheit	°

Für die Merkmalsausprägungen des Lautwertes werden segmentale Symbole herangezogen und die alphabetischen Zeichen

orale Kurzvokale	<i>o, ɔ, œ, e, æ, a, u, y, i</i>
nasale Kurzvokale	<i>õ, õ̃, œ̃, ě, æ̃, â, û, ĵ, ĭ</i>
orale Langvokale	<i>ō, ȯ, œ̄, ē, ǣ, ā, ū, ŷ, ĭ</i>
nasale Langvokale	<i>ȭ, ȭ̃, œ̄̃, ě̄, ǣ̃, ā̄, ū̄, ŷ̄, ĭ̄</i>
Halbvokale	<i>w, j</i>
Liquiden	<i>ɹ, R, L, ĩ, r, R, l, l</i>
Nasale	<i>ɱ, ɲ, ñ, m, n, ŋ</i>
Sibilanten	<i>s, z</i>
Frikative	<i>f, ɸ, β, χ, h, ʔ, v, ð, ʒ</i>
Okklusive	<i>p, t, k, b, d, g</i>

gebraucht. Soweit es die zur Verfügung stehenden Zeichensätze zuließen, wurden zur Angabe des Lautwertes dieselben Symbole ausgewählt, wie sie auch in der Arbeit von Cathey vorkommen. Bei Nasalen und Liquiden bezeichnet ein horizontaler Strich als Diakritikum Stimmlosigkeit, $\bar{r}$ bzw. $\bar{R}$ die Varianten des dem Runen- R zugeschriebenen Lautes, und die mit Punkt versehenen Symbole für Liquiden die dentalen gegenüber den punktlos dargestellten kakkuminalen Spielarten. Punktierte Symbole für Nasale werden bei Cathey im Vorspann aufgeführt, treten aber im weiteren Verlauf nirgendwo mehr in Erscheinung; ausweislich der binären Merkmale handelt es sich um velare n . Bei den Frikativen stehen f und $\bar{b}$ für die bilabialen, f und v für die labiodentalen Varianten; h gibt den reinen Hauchlaut wieder.

In allen Merkmalsdimensionen repräsentiert $\square$ die leere Merkmalsausprägung.

Mit dem eben eingeführten Notationssystem für lautliche Einheiten werden alle sonst nur implizit gegebenen Merkmalsausprägungen, wie die Unbetontheit von Konsonanten oder die Position im Wort, konsequent explizit gemacht. Die gleichfalls üblicherweise nicht gesondert notierte Aufeinanderfolge von lautlichen Einheiten wird durch den Verkettungsoperator $\nearrow$ gekennzeichnet.

Sofern nicht nur einzelne Punkte des Merkmalsraumes, wie bei den lautlichen Einheiten von Wortformen, sondern Teilmengen des Merkmalsraumes spezifiziert werden sollen, wie es für die Regeln des Lautwandels erforderlich ist, werden die gewünschten Merkmalsausprägungen, für jede Merkmalsdimension separat, nacheinander und durch Kommata getrennt innerhalb eines Klammerpaares $\langle \quad \rangle$ aufgelistet; die zu den unterschiedlichen Merkmalsdimensionen gehörenden Klammersausdrücke werden wieder durch $\oplus$ verknüpft. Der so entstehende viergliedrige Term ist aufzufassen als ein aus vier Mengen gebildetes kartesisches Produkt; es werden also alle möglichen Kombinationen aus den aufgeführten Merkmalsausprägungen betrachtet. Die Aufeinanderfolge von Mengen lautlicher Einheiten wird, ebenso wie oben, durch den Verkettungsoperator $\nearrow$ ausgedrückt.

Um dem Benutzer des Programmes einen schnellen ersten Überblick über die berechneten Lautwandelprozesse zu ermöglichen, gelangt neben der vierdimensionalen Langformdarstellung eine aus dieser gewonnene, den üblichen Gepflogenheiten entsprechende Kurzform zur Ausgabe, in der nur die Lautwerte aufgeführt sind und auf den Verkettungsoperator verzichtet wird.

Es folgt nun zur Illustration je ein Beispiel zur Schreibung einer Wortform (1) und einer Regel des Lautwandels (2).

(1) Zur Modellierung seiner Lautgeschichte wird das hypothetische Wort *tastjan* eingegeben als

$$\begin{aligned} \text{d)} \quad & \leftarrow \oplus \leftrightarrow \oplus^\circ \oplus t \cap \rightrightarrows \oplus \leftrightarrow \oplus !! \oplus a \cap \rightrightarrows \oplus \leftrightarrow \oplus^\circ \oplus s \cap \rightrightarrows \oplus \leftrightarrow \oplus^\circ \oplus t \cap \\ & \rightrightarrows \oplus \leftrightarrow \oplus^\circ \oplus j \cap \rightrightarrows \oplus \leftrightarrow \oplus^\circ \oplus a \cap \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus n \end{aligned}$$

In Kurzform lautet die Programmausgabe

$$\begin{aligned} & \textit{tastjan} \rightarrow 1\text{A} \rightarrow \textit{tastj\~{a}n} \rightarrow 6\text{B}, 3 \rightarrow \textit{t\ae stj\~{a}n} \rightarrow 11\text{C}, 1 \rightarrow \textit{t\ae stj\~{a}} \rightarrow 13\text{H}, 5 \rightarrow \textit{t\ae st\~{a}} \rightarrow \\ & 17\text{D} \rightarrow \textit{test\~{a}} \end{aligned}$$

Die etwas schwer lesbare Langform der Programmausgabe wird im untenstehenden Diagramm der besseren Übersichtlichkeit wegen abschnittsweise der Kurzform gegenübergestellt.

Die zwischen zwei Pfeilen eingeschlossenen Kombinationen aus Ziffern und Großbuchstaben stehen für die zur Anwendung gekommenen Regeln des Lautwandels und entsprechen der Numerierung in Catheys Arbeit.

(2) Üblicherweise werden Lautwandelregeln angegeben in der Form $Q \rightarrow Z/L_R$, wobei der Unterstrich $_$ die Transformationsstelle bezeichnet, die Quelle Q die Gestalt der Transformationsstelle vor Anwendung der Transformation, das Ziel Z die Gestalt der Transformationsstelle nach Anwendung der Transformation, und L bzw. R den Links- bzw. Rechtskontext. Alternativ und etwas weniger kompakt könnte man stattdessen, wie bei kontextsensitiven Grammatiken, auch $L \sqcap Q \sqcap R \rightarrow L \sqcap Z \sqcap R$ schreiben, wodurch das naheliegende Prinzip angedeutet ist, nach dem im Programm derartige Transformationsregeln abgearbeitet werden.

Den Ausfall von *þ* zwischen unmittelbar vorausgehendem Vokal und unmittelbar nachfolgenden Allophonen von */* unter Ersatzdehnung des Vokals, sofern er kurz ist, faßt Cathey unter Bezugnahme auf Heusler (1964: Par. 163.1) und Noreen (1923: Par. 236) zusammen als

Rule 2L: Loss of *þ* before *l*

$$\check{V} \not\models p \rightarrow \nabla \emptyset \quad / \quad [1]$$

Die Symbole für Laute und Lautgruppen sind relativ zum gerade betrachteten zeitlichen Schnitt zu interpretieren und explizit zu machen. Ein Teil von Regel 2L wurde in der Form

Bez	2L
Q	$\nabla \leftarrow, \sqsubset \nabla \oplus \nabla \rightsquigarrow, \leftrightarrow \nabla \oplus \nabla !! , !, {}^\circ \nabla \oplus \nabla a, x, e, i, o, q, \alpha, u, y, \tilde{a}, \tilde{x}, \tilde{e}, \tilde{i}, \tilde{o}, \tilde{\rho}, \tilde{\alpha}, \tilde{u}, \tilde{y} \nabla \sqcap$ $\nabla \sqsubset \nabla \oplus \nabla \leftrightarrow \nabla \oplus \nabla {}^\circ \nabla \oplus \nabla p \nabla$
Z	$\nabla \leftarrow, \sqsubset \nabla \oplus \nabla \rightsquigarrow, \leftrightarrow \nabla \oplus \nabla !! , !, {}^\circ \nabla \oplus \nabla \bar{a}, \bar{x}, \bar{e}, \bar{i}, \bar{o}, \bar{\rho}, \bar{\alpha}, \bar{u}, \bar{y}, \tilde{\bar{a}}, \tilde{\bar{x}}, \tilde{\bar{e}}, \tilde{\bar{i}}, \tilde{\bar{o}}, \tilde{\bar{\rho}}, \tilde{\bar{\alpha}}, \tilde{\bar{u}}, \tilde{\bar{y}} \nabla \sqcap$ $\nabla \square \nabla \oplus \nabla \square \nabla \oplus \nabla \square \nabla \oplus \nabla \square \nabla$
L	$\nabla \square \nabla \oplus \nabla \square \nabla \oplus \nabla \square \nabla \oplus \nabla \square \nabla$
R	$\nabla \rightarrow, \sqsubset \nabla \oplus \nabla \leftrightarrow \nabla \oplus \nabla {}^\circ \nabla \oplus \nabla l, \underline{l}, \bar{l}, \tilde{l} \nabla$
Beg	22
End	23
Prior	1
Richtg	reg
...	...

innerhalb einer Tabelle zur relativen Chronologie abgelegt, auf die das Programm zugreift und in der alle Lautwandelregeln in normierter Gestalt gesammelt sind. Jede Regel wird so formuliert, daß die Anzahl der zugehörigen Ver-

kettungen lautlicher Einheiten von Quelle und Ziel übereinstimmt und einander bei der Transformation entsprechende Verkettungen lautlicher Einheiten auf einander entsprechende Positionen zu liegen kommen; zusätzlich werden jeder Regel Informationen zu Wirkungszeitraum und Wirkungsdauer (Beg, End), zu Prioritätsstufe bei gleichzeitigem Auftreten mehrerer Regeln (Prior) und Angaben zur Durchlaufrichtung (Richtg mit den beiden Ausprägungen prog und reg für progressiv und regressiv) beigegeben. In einer festgelegten Reihenfolge werden dann alle möglichen Verkettungen von lautlichen Einheiten aus dem Regellinksbereich $L \cap Q \cap R$ und dem Regelrechtsbereich $L \cap Z \cap R$ gebildet; jedes Paar aus einander entsprechenden Verkettungen aus Regellinksbereich und Regelrechtsbereich wird als Termersetzung auf die aktuelle Gestalt angewandt, die ein Wort, dessen Geschichte modelliert werden soll, bei Eintritt des Wirkungszeitraumes der Regel aufweist.

Für das vorliegende Beispiel hat man in Quelle und Ziel jeweils $(2 \times 2 \times 3 \times 18) \times (1 \times 1 \times 1 \times 1) = 216$ Verkettungen, für den Linkskontext sind es $1 \times 1 \times 1 \times 1 = 1$, für den Rechtskontext $2 \times 1 \times 1 \times 4 = 8$ Verkettungen; somit enthalten Regellinksbereich und Regelrechtsbereich je $1 \times 216 \times 8 = 1728$ Verkettungen lautlicher Einheiten, also sind 1728 Termersetzungen zur Abarbeitung von Regel 2L vorzunehmen.

In abschnittsweiser Darstellung erhält man für den Regellinksbereich von $2L$

[illegible]

und in dem entsprechender Form für den Regelrechtsbereich

[illegible]

woraus man die Termersetzungen

$$\begin{aligned} & \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \leftarrow \oplus \rightsquigarrow \oplus !! \oplus a \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus p \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus l \star \rightarrow \\ & \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \leftarrow \oplus \rightsquigarrow \oplus !! \oplus \bar{a} \star \sqcap \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus l \star, \end{aligned}$$

$$\begin{aligned} & \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \leftarrow \oplus \rightsquigarrow \oplus !! \oplus \bar{a} \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus p \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus l \star \rightarrow \\ & \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \leftarrow \oplus \rightsquigarrow \oplus !! \oplus \bar{a} \star \sqcap \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus l \star, \end{aligned}$$

...

$$\begin{aligned} & \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus \bar{y} \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus p \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus \bar{l} \star \rightarrow \\ & \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus \bar{y} \star \sqcap \star \square \oplus \square \oplus \square \oplus \square \star \sqcap \star \rightarrow \oplus \leftrightarrow \oplus^\circ \oplus \bar{l} \star \end{aligned}$$

entnimmt.

Die Längen der Regelbereiche der relativen Chronologie vom Urgermanischen zum Altisländischen liegen zwischen 29 und 13847 Zeichen; die derzeit eingesetzte Programmierungsumgebung läßt bis zu 65535 Zeichen zu.

Um den Anschluß an die Fachliteratur zum Altisländischen zu gewährleisten, wird ein zweites Notationssystem für die altisländischen Zielformen eingeführt, das sich an die in Wörterbüchern üblichen Normalisierungen anlehnt. Das hierzu benutzte Alphabet wird beschrieben durch die nachstehende Übersicht.

Kurzvokale	<i>i, y, u, e, ø, o, ø, a</i>
Langvokale	<i>í, ý, ú, é, æ, ó, æ, á</i>
Halbvokale	<i>v, j</i>
Liquiden	<i>r, l</i>
Nasale	<i>m, n</i>
Sibilanten	<i>s</i>
Frikative	<i>f, þ, h, ð</i>
Okklusive	<i>p, t, k, b, d, g</i>
Affrikaten	<i>z, x</i>

Positionsbedingt bezeichnen die Symbole *b, g* für stimmhafte Okklusive auch die Lautung der ihnen entsprechenden Frikative. Bei *v* differieren die Ansichten, wann eine Zuordnung zu den Halbvokalen oder zu den Frikativen zu treffen ist.

Zur Umwandlung in altisländische Notation werden die Kurzformen der Endprodukte der Modellierung des Lautwandels einigen Transformationen unterzogen, die in der anschließend gegebenen Reihenfolge vorgenommen werden.

(a) Kurzvokale

$\bar{a} \rightarrow a, \bar{e} \rightarrow e, \bar{i} \rightarrow í, \bar{o} \rightarrow o, \bar{u} \rightarrow u, \bar{y} \rightarrow y, \bar{x} \rightarrow e, \bar{\tilde{x}} \rightarrow e, \bar{\alpha} \rightarrow \emptyset, \bar{\tilde{\alpha}} \rightarrow \emptyset, \bar{\varphi} \rightarrow \varphi.$

(b) Langvokale

$\bar{a} \rightarrow \acute{a}, \bar{\tilde{a}} \rightarrow \acute{a}, \bar{e} \rightarrow \acute{e}, \bar{\tilde{e}} \rightarrow \acute{e}, \bar{i} \rightarrow \acute{i}, \bar{\tilde{i}} \rightarrow \acute{i}, \bar{o} \rightarrow \acute{o}, \bar{\tilde{o}} \rightarrow \acute{o}, \bar{u} \rightarrow \acute{u}, \bar{\tilde{u}} \rightarrow \acute{u},$
 $\bar{y} \rightarrow \acute{y}, \bar{\tilde{y}} \rightarrow \acute{y}, \bar{x} \rightarrow \acute{x}, \bar{\tilde{x}} \rightarrow \acute{x}, \bar{\alpha} \rightarrow \acute{\alpha}, \bar{\tilde{\alpha}} \rightarrow \acute{\alpha}, \bar{\varphi} \rightarrow \acute{a}, \bar{\tilde{\varphi}} \rightarrow \acute{a}$

(c) Konsonanten

$\bar{l} \rightarrow l, \bar{l} \rightarrow l, \bar{\tilde{l}} \rightarrow l, \bar{n} \rightarrow n, \bar{m} \rightarrow m, \bar{r} \rightarrow r, \bar{z} \rightarrow g, \bar{f} \rightarrow p, \bar{ks} \rightarrow x, \bar{vn} \rightarrow fn.$

Möglicherweise sind in einer zukünftigen Version des Programmes oder bei Erweiterung der Stichprobe von Kandidaten zur Modellierung noch Ersetzungen zur Bildung von Affrikaten hinzuzufügen.

Wie man leicht feststellt, vertreten gleiche Zeichen in den unterschiedlichen Darstellungen unterschiedliche Lautungen; so stehen beispielsweise $\bar{x}$ und $\bar{\alpha}$ in der für die Modellierung verwendeten Notation für Kurzvokale, beim Altisländischen dagegen für Langvokale.

4.2. Die Benutzeroberfläche

[illegible]

Modellierung der Entwicklung von aisl. *dýr* 'wildes Tier' aus dem urgerm. Ansatz **deuzā* (Cathey 1967:36)

4.3. Beschreibung des Leistungsumfangs der Modellierung

Wenn man zur Ermittlung der Trefferquote strenge Kriterien zugrundelegt und gewisse Abweichungen selbst in nur einem Merkmal pro Zeichen als fehlerhaft wertet, wie in der Übersicht

berechnete Form	Zielform
<i>deildā</i>	<i>deilda</i>
<i>bāss</i>	<i>báss</i>
<i>fæddā</i>	<i>fædda</i>
<i>lityt</i>	<i>litit</i>
<i>letr</i>	<i>leðr</i>
<i>þorv</i>	<i>þorf</i>

durch Hinterlegung angedeutet, wobei unterschiedliche schreiberische Gepflogenheiten für die berechneten Formen und die Zielformen zu beachten sind (vgl. die obenstehenden Erläuterungen zur Umwandlung in altisländische Notation), so erhält man bei insgesamt 394 Lemmata 218 Treffer, also eine Trefferquote von 55,3%.

Exemplarisch werden einige lange, als richtig zu bewertende Lautgeschichten angeführt.

Urgerm. Ansatz	Lautgeschichte	Zielform
<i>fawizē</i>	<i>fawizē</i> → 3A → <i>fawirē</i> → 6B,2 → <i>fæwirē</i> → 11A,1 → <i>fæwrē</i> → 13D → <i>fæwrī</i> → 15C → <i>fæwrī</i> → 16A,1 → <i>færi</i> → 18B,1 → <i>færi</i>	<i>færi</i>
<i>kunþaz</i>	<i>kunþaz</i> → 1A → <i>kunþaz</i> → 3A → <i>kunþar</i> → 3C → <i>kunðar</i> → 4F → <i>kunnar</i> → 11A,1 → <i>kunnr</i> → 15C → <i>kunnr</i> → 17B → <i>kūðr</i> → 18D → <i>kūðr</i>	<i>kuðr</i>
<i>lazjan</i>	<i>lazjan</i> → 1A → <i>lazjān</i> → 3B → <i>lazzjān</i> → 4C,1 → <i>lagzjān</i> → 4C,2 → <i>laggjān</i> → 6B,3 → <i>læggjān</i> → 11C,1 → <i>læggjā</i> → 17D → <i>leggjā</i>	<i>leggja</i>
<i>maþlā</i>	<i>maþlā</i> → 1E → <i>maþlā</i> → 2L → <i>mālā</i> → 10A → <i>māl</i>	<i>māl</i>
<i>nēχwizē</i>	<i>nēχwizē</i> → 2H → <i>nāχwizē</i> → 3A → <i>nāχwirē</i> → 4K,1 → <i>nāχirē</i> → 6B,2 → <i>nāχirē</i> → 6C → <i>nāχirē</i> → 11A,1 → <i>nāhrē</i> → 1E → <i>nāhrē</i> → 13A,1 → <i>nāhrē</i> → 13D → <i>nāRī</i> → 15C → <i>nāri</i>	<i>næri</i>

Urgerm. Ansatz	Lautgeschichte	Ziel-form
<i>sankwjan</i>	<i>sankwjan</i> → 1A → <i>sānkwjān</i> → 5A → <i>sānkwjān</i> → 6A → <i>sākkwjān</i> → 6B,3.5 → <i>sækkwjān</i> → 7D,2 → <i>sœkkwjān</i> → 11C,1 → <i>sœkkwjā</i> → 13H,6 → <i>sœkkwā</i> → 17A → <i>sœkkbā</i> → 18B,1 → <i>sœkkvā</i>	<i>søkkva</i>

4.4. Auswertung der Abweichungen

Es folgt nun eine Auswahl an fehlerhaften Berechnungen einschließlich der wesentlichen Abweichungen. Die bei den fehlerhaften Berechnungen zu beobachtenden Abweichungen sind auch im nächsten Kapitel, bei der Formen-erzeugung schwacher verbaler Simplizia, wieder anzutreffen.

Abkür-zung	Beschreibung der Abweichung
assim	Geminatenkürzung, Dreikonsonantenregel oder Assimilations-erscheinungen unter Beteiligung von Konsonanten
<i>a</i> -Uml	Senkung von <i>u</i> oder <i>a</i> -Umlaut
bre-qm	Brechung oder Quantitätenmetathese
<i>v</i> -Probl	Assimilationserscheinungen unter Beteiligung von <i>v</i> bzw. <i>w</i>

Urgerm. Ansatz	Lautgeschichte	Ziel-form	Abwei-chung
<i>axsulaz</i>	<i>axsulaz</i> → 2J,1 → <i>aksulaz</i> → 3A → <i>aksulaR</i> → 9A,2 → <i>øksulaR</i> → 11A,1 → <i>øksulR</i> → 15C → <i>øksulr</i>	<i>øxull</i>	assim
<i>þunxtō</i>	<i>þunxtō</i> → 1A → <i>þūnxtō</i> → 4I → <i>þūnttō</i> → 5A → <i>þūnttō</i> → 6A → <i>þūttō</i> → 12B → <i>þūttā</i> → 18A,1 → <i>þūttā</i>	<i>þóttā</i>	<i>a</i> -Uml
<i>wulfaz</i>	<i>wulfaz</i> → 3A → <i>wulfaR</i> → 3C → <i>wulbaR</i> → 4H,2 → <i>wolbaR</i> → 10D,4 → <i>olbaR</i> → 11A,1 → <i>olbR</i> → 15B → <i>ōlbR</i> → 15C → <i>ōlbr</i> → 18B,1 → <i>ōlvr</i>	<i>ulfr</i>	<i>a</i> -Uml
<i>χrunkwō</i>	<i>χrunkwō</i> → 1A → <i>χrūnkwō</i> → 1E → <i>χrūnkwō</i> → 1I → <i>hrūnkwō</i> → 4K,8 → <i>hrūnkō</i> → 5A → <i>hrūpkō</i> → 6A → <i>hrūkkō</i> → 7B,2 → <i>hrōkkō</i> → 12B → <i>hrōkkā</i>	<i>hrukka</i>	<i>a</i> -Uml
<i>fuzlaz</i>	<i>fuzlaz</i> → 3A → <i>fuzlaR</i> → 4H,2 → <i>fozlaR</i> → 11A,1 → <i>fozlr</i> → 15C → <i>fozlr</i> → 18B,1 → <i>fozlr</i>	<i>fugl</i>	assim, <i>a</i> -Uml

Urgerm. Ansatz	Lautgeschichte	Ziel-form	Abwei-chung
<i>etunaz</i>	<i>etunaz</i> → 1A → <i>etūnaz</i> → 3A → <i>etūnar</i> → 9A,1 → <i>ætūnar</i> → 11A,1 → <i>ætūnr</i> → 15C → <i>ætūnr</i> → 16D,1 → <i>ætūnn</i> → 18A,3 → <i>ætūn</i>	<i>jōtunn</i>	assim, bre-qm
<i>seχwan</i>	<i>seχwan</i> → 1A → <i>seχwān</i> → 3F → <i>sēχwān</i> → 4K,1 → <i>sēχān</i> → 6C → <i>sēhān</i> → 11C,1 → <i>sēhā</i> → 13A,1 → <i>sēā</i>	<i>sjá</i>	bre-qm
<i>χnewwaz</i>	<i>χnewwaz</i> → 1E → <i>χnewwaz</i> → 1I → <i>hnewwaz</i> → 2N → <i>hnezwaz</i> → 3A → <i>hnezwaz</i> → 3B → <i>hnezwaz</i> → 4C,1 → <i>hnezwaz</i> → 4C,2 → <i>hneggwar</i> → 4K,5.33 → <i>hneggaz</i> → 8A,3 → <i>hneaggaz</i> → 11A,1 → <i>hneaggr</i> → 15A,1 → <i>hnjaggr</i> → 15C → <i>hnjaggr</i>	<i>hnøggr</i>	bre-qm, v-Probl
<i>arχwōz</i>	<i>arχwōz</i> → 3A → <i>arχwōr</i> → 4K,5.33 → <i>arχōr</i> → 5A → <i>arχōr</i> → 6C → <i>aṛhōr</i> → 12B → <i>aṛhar</i> → 13A,1 → <i>aṛar</i> → 15C → <i>aṛar</i>	<i>orvar</i>	v-Probl
<i>χawwan</i>	<i>χawwan</i> → 1A → <i>χawwān</i> → 1I → <i>hawwān</i> → 2N → <i>hazwān</i> → 3B → <i>hazwān</i> → 4C,1 → <i>hagzwān</i> → 4C,2 → <i>haggwān</i> → 4K,5.33 → <i>haggān</i> → 11C,1 → <i>haggā</i>	<i>hoggva</i>	v-Probl

Die durch Regel 18A,3 hervorgerufene unerwünschte Kürzung von Konsonantengruppen ist innerhalb der Gesamtheit der falsch berechneten Wortgeschichten 18 mal die Ursache für Fehler; bei den richtig berechneten Wortgeschichten kommt sie nirgends zur Anwendung. Daher ist über eine Streichung von Regel 18A,3 nachzudenken.

5. Wortformenerzeugung für altisländische schwache verbale Simplicia

5.1. Aufbau und Zielsetzung

Ein Katalogisierungsansatz ist eine sehr erfolgversprechende und naheliegende Möglichkeit, für die Elemente einer von grammatischen Gesichtspunkten aus als zusammengehörig anzusehenden Menge von Grundformen Paradigmen in großer Zahl automatisch zu erzeugen. Er besteht darin, dem Beispiel der Grammatiken zu folgen und einem Musterparadigma eine Liste von Wörtern zur Seite zu stellen, die in gleicher Weise flektieren. Hierzu ist die gegebene Menge von Grundformen geeignet in Klassen einzuteilen und zu jeder Klasse ein prototypisches Flexionsschema zu erstellen, das neben den für die gesamte Klasse charakteristischen konstanten Zeichenketten auch Variablen enthält, die

durch Teilzeichenketten der einzelnen Repräsentanten der Klasse besetzt werden. Dieses Vorgehen führt im Ergebnis also zu einer Katalogisierung des gewählten Ausschnittes des Wortschatzes in der Weise, daß jeder Grundform das zu ihr passende Flexionsschema zugeordnet wird, wodurch eine vollständige Bereitstellung fehlerfreier Paradigmen erreicht werden kann. Trotz dieses unbestreitbaren Vorteils ist zu wissenschaftlichen Zwecken das eben skizzierte Vorgehen alleine sowie ein daraus hervorgehendes Datenverarbeitungsprogramm für unbefriedigend zu halten, da die für die Klasseneinteilung zu berücksichtigende, oftmals umfangreiche sprachhistorische Information schon vom Ansatz her nirgends sichtbar gemacht werden kann, und da etwa zusätzliche betrachtenswerte Grundformen, auch vielleicht solche virtueller Art, nicht ohne weiteren Aufwand in ein derartiges Programm integriert werden können.

Der eben genannten Nachteile wegen wird, ganz im Sinne der eingangs vorgestellten Konzeption, dem Algorithmus zur Wortformenerzeugung ein diachroner Ansatz zugrundegelegt. Die Eingabedaten sind Zeichenketten über einem zur normalisierten Schreibung einer früheren Sprachstufe verwendeten Alphabet, also Teilzeichenketten von Transponaten, Rekonstrukten oder bezeugten Wörtern, auf die im Programmverlauf Regeln zur Formenbildung und zum Lautwandel wirken, um als Ausgabedaten schließlich Paradigmen der anvisierten späteren Zeitstufe zu erhalten.

Unter den in das vorliegende Formenerzeugungsprogramm eingearbeiteten Aspekten stehen morphologische Gesichtspunkte im Vordergrund. Für die Formenerzeugung altisländischer schwacher verbaler Simplizia werden urgermanische Rekonstrukte einer Zeitstufe herangezogen, in der die Grimmsche Lautverschiebung, das Vernersche Gesetz und das Festwerden der Wortbetonung als bereits vollzogen gelten. Das Dentalpräteritum wird modelliert durch Komposition mit den entsprechenden Formen des Rekonstruktes zum Verbum *tun*, wie das auch von manchen sprachwissenschaftlichen Ansätzen, z.B. (Lühr 1984:41ff.), unter dem Begriff der Kompositionstheorie nahegelegt wird, und die verbale Stammbildung wird durch Anfügen nichtredundanter Teile der urgermanischen Stammbildungsmorpheme vollzogen. Da keine Katalogisierung im oben diskutierten Sinne betrieben werden soll, wird auf die in den Lehrbuchgrammatiken übliche Durchnummerierung der altisländischen schwachen Verbalklassen als Eingabedatum konsequenterweise verzichtet; stattdessen werden die die Flexion bestimmenden morphologischen Parameter herangezogen. Somit muß, zumindest bei den Verben mit *jan* als Stammbildungsmorphem, als Diskriminante zur Unterscheidung der Flexionstypen auf das

Wurzelgewicht zurückgegriffen werden, dessen Festlegung sich immer dann schwierig gestaltet, wenn man es zusätzlich mit Konsonantengemination oder Verschärfung zu tun hat oder in irgendeinem Abschnitt des Beobachtungszeitraums mit leeren Wurzelabglitten konfrontiert wird. Die für die eben genannten Aufgaben bestimmten Unterprogramme des Wortformenerzeugungsalgorithmus müssen also sehr subtil aufeinander abgestimmt werden, was sich als eine der komplexeren Problematiken bei der Programmentwicklung erwiesen hat.

Um überparadigmatischen Ablaut, der durch frühurgermanische Lautvorgänge verdunkelt wird, explizit zu machen, wird für den Silbengipfel des lexikalischen Morphems auch vorgermanischer Vokalismus zugelassen und eine versuchsweise Einordnung hinsichtlich Ablautstufe und Ablautreihe vorgenommen.

Zu den nicht primär morphologischen Aspekten bei der Formenerzeugung zählen paradigmatische Einflüsse auf die Wortformen. Obwohl sich viele Analogien in ähnlicher Weise als Transformationen darstellen lassen wie Regeln des Lautwandels, wenn man neben Lautungen auch grammatische Termini als formale Größen auffaßt, finden Analogien bislang nur spärliche Aufmerksamkeit und werden im Formenerzeugungsprogramm auch noch nicht als Transformationen verarbeitet. Begründet ist dies dadurch, daß Analogien einer ausführlichen Analyse der an ihr teilnehmenden Elemente bedürfen, will man nicht der Beliebigkeit Tür und Tor öffnen, und diese Arbeit ist eben nicht im erforderlichen Umfange geleistet.

Die Erzeugung sogenannter Mischparadigmen bleibt vorbehalten; ebenso sind Defektivitäten beliebiger Art und Ursache in den Paradigmen von der Verarbeitung derzeit ausgeschlossen.

Mit dem vorgestellten Formenerzeugungsprogramm kann man mehrere weitergehende Ziele verfolgen. Ein Nutzen, der sich aus jeglichem Wortformenerzeugungsprogramm ergibt, ist die Möglichkeit, die durch Handschriften oder über normalisierte Textausgaben tradierte sprachliche Realität mit der in den Formenlehren der Grammatiken geleisteten Abstraktion zu vergleichen und ggf. Ungereimtheiten oder Lücken in dieser Abstraktion aufzuzeigen. Der diachrone Ansatz erlaubt es zusätzlich, der sprachlichen Realität die Lautlehren der Grammatiken gegenüberzustellen, um so zu verfeinerten Regeln des Lautwandels und zu stimmigen Aufstellungen der bis jetzt meist nur oberflächlich behandelten relativen Chronologien zu gelangen.