



universität
wien

DIPLOMARBEIT

Titel der Diplomarbeit

„Reime und der Zugriff aufs Lexikon“

Verfasserin

Michaela Rausch

angestrebter akademischer Grad

Magistra der Philosophie (Mag. phil.)

Wien, 2013

Studienkennzahl lt. Studienblatt:

Studienrichtung lt. Studienblatt:

Betreuerin:

A 328

Allgemeine/Angewandte Sprachwissenschaft

Ao. Univ.-Prof. Dr. Chris Schaner-Wolles

Inhaltsverzeichnis

1. Einleitung	1
1.1. Forschungsfragen und Hypothesen	1
1.2. Aufbau	2
2. Was sind Reime?	3
2.1. Reimtypen	3
2.2. Die Silbe	5
2.2.1. Bestandteile der Silbe	6
2.2.1.1. Onset	6
2.2.1.2. Nukleus	8
2.2.1.3. Koda	9
2.2.1.4. Die phonologische Komponente „Reim“	10
2.2.2. Die Sonoritätshierarchie	13
2.2.3. Silbengrenzen	15
2.3. Mehrsilbige Reime	16
3. Überblick über den aktuellen Forschungsstand zu Reimen in der Psycho-, Patho- und Neurolinguistik	19
3.1. Silben und Reime im Spracherwerb	19
3.2. Die phonologische Bewusstheit	21
3.3. Reime bei sprachgesunden Erwachsenen	25
3.4. Reime bei aphasischen Erwachsenen	30
3.5. Neurologische Verarbeitung von Reimen	38
4. Theoretischer Hintergrund	45
4.1. Das mentale Lexikon	45
4.2. Der Zugriff aufs Lexikon	49
4.3. Sprachproduktionsmodelle	51
4.3.1. Das Modell von Caramazza	54
4.3.2. Serielle Sprachproduktionsmodelle	57
4.3.3. Das Sprachproduktionsmodell von Levelt et al. (1999)	58
4.3.4. Interaktive Sprachproduktionsmodelle	62
4.3.5. Das Sprachproduktionsmodell von Dell (1986)	63
4.3.6. Die schriftliche Sprachproduktion	68

4.4. Zusammenfassung	74
5. Experiment	75
5.1. Priming-Experimente	75
5.2. Das Experiment von Rapp und Samuel (2002)	78
5.2.1. Aufbau	78
5.2.2. Methode	79
5.2.3. Ergebnisse	81
5.3. Adaption des Experiments für das Deutsche	81
5.4. Aufbau des Experiments	82
5.5. Methode	84
5.5.1. Probanden	84
5.5.2. Material und Design	84
5.5.3. Verfahren	87
5.6. Ergebnisse	88
6. Diskussion und Zusammenfassung	89
7. Bibliographie	93
8. Abbildungsverzeichnis	103
9. Anhang	105
A. Material (Experiment)	105
A.1. Items mit einsilbigen Reimen	105
A.2. Items mit zweisilbigen Reimen	108
A.3. Füllitems	111
A.4. Übungsite ms	114
A.5. Beschreibung	114
B. Zusammenfassung	116
C. Lebenslauf	117

1. Einleitung

Reime begleiten den Menschen im Alltag seit seiner Kindheit. Angefangen von Kinderliedern und Abzählreimen, über Gedichte und Redewendungen, ja bis hin zu Werbeslogans und Gedächtnisstützen - Reime sind überall zu finden. Doch was macht sie genau aus?

In der vorliegenden Arbeit wird versucht, Reime von verschiedenen Seiten zu beleuchten, um vielleicht ihre Wirkung und Eigenheiten als Ganzes besser erkennen zu lassen. Was sind Reime? Warum halten wir sie für schön und ganzheitlich, was macht sie prominent für uns? Und was kann mit ihnen untersucht werden, was können sie über die Sprachverarbeitung des Menschen aussagen? Dies und weitere Aspekte werden in den nachfolgenden Kapiteln behandelt.

Ich habe dieses Thema gewählt, da ich verschiedene Teile der Sprachverarbeitung anschneiden wollte: Phonologie und phonologische Bewusstheit, lexikalischer Wortabruf und Schriftsprache. Dies war mir in Zusammenhang mit dem Thema „Reime“ möglich, da es in all diesen Gebieten der wissenschaftlichen Literatur behandelt wird. Ich halte Reime für sehr interessant, da sie sowohl in unserem Alltag, als auch in der Wissenschaft Fuß gefasst haben.

1.1. Forschungsfragen und Hypothesen

In den meisten Studien wird mit einsilbigen Reimen gearbeitet, meistens nur im Wortkontext. Ich möchte gerne ein- und mehrsilbige Reime in einem größeren Kontext, der Satzebene, bei sprachgesunden Erwachsenen untersuchen. Dazu soll das Experiment von Rapp und Samuel (2002) repliziert und fürs Deutsche adaptiert werden. Die Autoren haben herausgefunden, dass Reim-Primes im Englischen den Abruf von sich reimenden Wörtern bei sprachgesunden Erwachsenen erhöhen (Näheres unter Kap 5 *Experiment*), dass somit semantische *und* phonologische Faktoren den Wortabruf beeinflussen. Es wird untersucht, ob die Ergebnisse von Rapp und Samuel

mit denen in dieser Arbeit vergleichbar sind, oder ob es möglicherweise sprachspezifische Unterschiede gibt. In dieser Arbeit werden drei Hypothesen überprüft:

- Hypothese 1 besagt, dass im Rahmen dieser Arbeit bei Sätzen mit einsilbigen Reimen vergleichbare Ergebnisse auftreten wie in der Studie von Rapp und Samuel (2002).
- Das Experiment von Rapp und Samuel (2002) wird auf zweisilbige Reime erweitert, um zu überprüfen, ob die Ergebnisse auf zweisilbige Wörter übertragbar sind. Hypothese 2 besagt, dass hier wieder eine erhöhte Reimantwort für Reim-Primes auftritt, möglicherweise aber in eingeschränkterem Umfang, da die Verarbeitung zweisilbiger Wörter mehr Arbeitsaufwand darstellt als die einsilbiger Wörter.
- Hypothese 3 besagt, dass die Probanden aufgrund der umfangreicheren Antwortmöglichkeit in der deutschen Adaption mehr unpassende Antworten (nicht nur bestehend aus Artikel und Nomen) geben werden. Mehr dazu unter Kapitel 5 *Experiment*.

1.2. Aufbau

Zunächst möchte ich in der Arbeit auf den Aufbau von Reimen eingehen. Es sollen Fragen beantwortet werden, wie: Was sind Reime - phonologisch gesehen, poetisch gesehen? Was umfasst ein einsilbiger, was ein mehrsilbiger Reim? Dann möchte ich einen Überblick über den aktuellen Forschungsstand mit Reimen geben - ab wann können Kinder reimen? Es sollen Studien vorgestellt werden, in denen Untersuchungen mit Reimen durchgeführt wurden. Die Verarbeitung von Reimen bei sprachgesunden Erwachsenen, sowie bei aphasischen Erwachsenen wird von den nächsten Unterkapiteln beleuchtet. Wie werden Reime im Gehirn verarbeitet? Dann wird der lexikalische Wortabruf besprochen, der dafür verantwortlich ist, welches Wort eine Person sagt, wenn sie an etwas denkt bzw. etwas Bestimmtes meint. Wovon wird der lexikalische Wortabruf beeinflusst, rein von der Semantik, oder auch von anderen Faktoren? Zu guter Letzt werden nach einer Einleitung über Priming-Experimente das Experiment für diese Arbeit und seine Ergebnisse präsentiert. Im Anhang befinden sich die Materialien dazu. Es folgt eine Zusammenfassung der in dieser Arbeit gewonnenen Erkenntnisse.

2. Was sind Reime?

Der Begriff „Reim“ stammt aus dem Französischen: *Rimer* bedeutet in etwa so viel wie „in Reihen anordnen“. Bereits im Mittelhochdeutschen existierte das Wort *ríme/-rim* (Bußmann ⁴2008, Glück ⁴2010).

Bußmann (⁴2008) definiert Reime als „sprachliches Kunstmittel in vielen Sprachen der Welt: Gleichklang zweier oder mehrerer Silben vom letzten betonten Vokal der Hauptsilbe [...] an“. Glück (⁴2010) beschreibt Reime als „Gleichklang von Wörtern, meist vom letzten betonten Vokal eines Wortstamms an“.

Es gibt verschiedene Arten von Reimen, die aufgrund ihrer Erscheinungsform oder ihres Erscheinungsortes unterschiedliche Bezeichnungen haben. Um einen Überblick über die Begrifflichkeiten und Unterteilungen zu geben, möchte ich diese verschiedenen Typen nun kurz vorstellen.

2.1. Reimtypen

Die klassische Form des Reimes ist der als Versverschluss dienende sogenannte **Endreim**. Er gilt als „häufigstes Kennzeichen von Gedichten“ (Glück ⁴2010, Stichwort *Reim*, S. 557) und die sich reimenden Wörter stehen am Ende der Verse oder Zeilen. Maas (2006, S.88) schreibt, dass der „Endreim [...] üblicherweise allein im Blick ist, wenn von Reim die Rede ist“.

Die Abfolge der Endreime in der Literatur und in Gedichten erfolgt anhand eines *Reimschemas*, das folgende Formen annehmen kann (Glück ⁴2010):

- Kreuzreim a-b-a-b
- Paarreim a-a-b-b
- Schweifreim a-a-b-c-c-b
- umarmender Reim a-b-b-a

2. Was sind Reime?

- verschränkter Reim a-b-c[d] a-b-c[d]

Andere Reimtypen sind der **Anfangsreim**, bei dem die ersten Wörter zweier oder mehrerer Verse reimen (vgl. *Krieg! ist das Losungswort, Sieg! und so klingt es fort ...*; Goethe, Faust II 9837/8, in: Glück ⁴2010), der **Binnen- oder Innenreim**, bei dem die sich reimenden Wörter innerhalb eines Verses befinden (vgl. *Sie blüht und glüht und leuchtet*, Heinrich Heine, in: Glück ⁴2010), sogenannte **Augenreime**, die nur das gleiche Schriftbild besitzen, sich aber phonologisch nicht reimen (vgl. engl. *cough* [ˈkʰɔ:f] „Husten“ vs. *tough* [ˈtʰʌf] „zäh“, Bußmann ⁴2008, Stichwort *Reim*, S. 558) und **Schüttelreime**, die dadurch entstehen, dass man durch Metathese die Anlaute der Reimwörter vertauscht, zum Beispiel: *Man bringt bei einem tollen Rausch/ sein Ich zu manchem Rollentausch* (Latzel, in: Bußmann ⁴2008).

Ein weiterer Typus ist der **Stabreim**, der eine andere Bezeichnung für die *Alliteration* darstellt. Bei der Alliteration handelt es sich um ein Stilmittel, bei dem der Onset bzw. der erste Laut isoliert wird und in einem zweiten Wort noch einmal als Onset auftritt, beispielweise in *Mann und Maus, Kind und Kegel, Haus und Hof* (Maas 2006, S. 118). Dieser Reimtyp ist der einzige Typus, der sich nicht zwangsläufig reimt, und gehört somit nicht in die Gruppe der engeren Definition von Reimen.

Reime können auch in **männliche** und **weibliche Reime** eingeteilt werden. Von einem männlichen oder auch „stumpfen“ Reim spricht man, wenn es sich um einen einsilbigen Reim handelt, der mit der Hebung einer betonten Silbe endet. Der weibliche oder auch „klingende“ Reim besteht aus einem zweisilbigen Reim, der mit einer Senkung einer unbetonten Silbe endet (Glück ⁴2010, Stichwort *Reim*, S. 557).

Ein weiteres Kapitel im Zusammenhang mit Reimen bildet die Unterscheidung in **reine** und **unreine Reime**. Bei einem unreinen Reim (oder *Halbreim*) handelt es sich um zwei oder mehr Wörter, die „keine vollständige Übereinstimmung in Vokal und Konsonanten“ aufweisen, zum Beispiel *flieht-blüht, Menschen-Wünsche* (Glück ⁴2010, Stichwort *Reim*, S. 557), oder *untertänig-König* (Wiese ²2006, S. 45). Im Gegensatz dazu kennzeichnen sich reine Reime dadurch aus, dass „alle Laute ab dem Nukleus der letzten betonten Silbe identisch sind“ (Altmann und Ziegenhain 2002, S. 85).

Der Endreim und die Alliteration sind als poetische Formen bei einsilbigen Wörtern wichtiger als die Assonanz (bei der der Vokal zweier oder mehrerer Wörter identisch ist - *send* und *bell*), die Konsonanz (bei der zwei Wörter einen Gleichklang bei unter-

schiedlichem Schriftbild besitzen - *send* und *hand*), als der „umgedrehte Reim“ (bei dem der Onset und der Nukleus übereinstimmen - *send* und *sell*) und der „Parareim“ (bei dem es sich um einen unreinen Reim handelt, der dadurch gekennzeichnet ist, dass die Vokale der beiden Wörter unterschiedlich sind bei gleichbleibender konsonantischer Umgebung - *send* und *sound*) (Duncan et al. 2007, S. 199).

Der Reimtyp, auf dem in dieser Arbeit das Augenmerk liegen soll, ist der klassische, reine Endreim, der aus einer oder mehreren Silben besteht. Dazu ist es wichtig, den Aufbau von Wörtern und die Beschaffenheit einer Silbe bzw. der deutschen Silben zu verstehen, weshalb ich nun näher darauf eingehen möchte.

2.2. Die Silbe

Eine ganz eindeutige Definition der Silbe gibt es nicht, da es sich um eine nicht einheitlich fassbare und identifizierbare sprachliche Einheit handelt, wie zum Beispiel ein Phonem. Sie besteht aus mehreren Elementen und besitzt verschwommene Grenzen. Im Folgenden werden verschiedene Definitionen vorgestellt, um ein Gesamtbild über die Silbe zu erlangen.

Der Begriff der Silbe stammt von dem griechischen *syl-labé* ab, was „das (beim Sprechen) Zusammengefasste“ (Bußmann ⁴2008, Stichwort *Silbe*, S. 624) und/oder „Fessel“ (Glück ⁴2010, Stichwort *Silbe*, S. 615) bedeutet. Bei der Silbe handelt es sich um die „kleinste suprasegmentale, d.h. lautübergreifende Einheit“ (Glück ⁴2010, S. 615). Bußmann (⁴2008) definiert die Silbe wie folgt:

[Die Silbe ist die] phonetisch-phonologische Grundeinheit des Wortes bzw. der Rede, die zwar intuitiv nachweisbar ist, wissenschaftlich aber nicht einheitlich definiert wird.

Wiese (²2006, S.34) schreibt über die Silbe:

It is a phonological unit organized around a syllabic peak (typically consisting of a vowel) of which speakers have a relatively high awareness.

Damit werden bereits einige wichtige Eigenschaften der Silbe angesprochen, die im weiteren Verlauf der Arbeit wieder aufgegriffen werden sollen: beispielsweise die intuitive Wahrnehmbarkeit der Silbe (siehe Kapitel 3.1 *Silben und Reime im Sprachewerb*), oder der Vokal als das „Zentrum der Silbe“ (siehe Unterkapitel 2.2.1.2 *Nukleus*).

Staffeldt (2010, S.133) beschreibt die Silbe als „die kleinste prosodische Einheit [...] und sie befindet sich als sprachliche Einheit zwischen den Phonemen und den vollständigen Wörtern“. Der Duden fügt dem hinzu, dass „Wortformen [...] also nicht direkt als Folgen von Lauten beschrieben [werden], sondern als Folgen von Silben“ (⁸2009, S. 37), wobei auch viele phonologische Regularitäten zum Großteil oft auf der Silbe beruhen (Wiese ²2006, S. 34). Somit kann man sagen, dass Wörter aus Silben und Silben aus Phonemfolgen bestehen. Laut Maas (2006, S. 115) ist die Silbe aber „nicht einfach eine lineare Kette aus ihren Segmenten“, was im folgenden Kapitel (Kapitel 2.2.1 *Bestandteile der Silbe*) näher erläutert werden soll.

Die Silbe gliedert die Sprache rhythmisch (Noack 2010), sie kann betont und unbetont sein, und trägt den Akzent (Duden ⁸2009). Im Folgenden werden weitere Punkte beschrieben, die im Allgemeinen aber nur für betonbare Silben gelten.

2.2.1. Bestandteile der Silbe

Eine Silbe besteht aus verschiedenen Komponenten, die jeweils mehrere Bezeichnungen haben. In dieser Arbeit werde ich allerdings die Begriffe **Onset**, **Nukleus** und **Koda** benutzen. Im Folgenden wird jede Komponente mit ihren Eigenschaften und Beschränkungen einzeln vorgestellt.

2.2.1.1. Onset

Der Onset heißt auch *Anfangsrand*, *Ansatz*, *Anlaut*, *Silbenanlaut* oder *Silbenkopf*, und bildet zusammen mit der Koda den *Silbenrand*. Er hat unterschiedliche Erscheinungsformen (Staffeldt 2010, S. 134):

- Wenn der Onset besetzt ist, spricht man von einer *gedeckten Silbe*.
 - Sie kann *einfach* sein (und besteht aus nur einem Laut), oder
 - *komplex* (und besteht aus mehreren Lauten).
- Wenn der Onset unbesetzt ist, handelt es sich um eine *nackte Silbe*.

Zweitere sind im Deutschen nicht sehr häufig. Sie treten nur als nicht betonbare Silben, somit als „zweite Silben“ auf, wie im Wort *gehen* ['ge:ən]. Fängt ein Wort im Schriftbild mit einem Vokal an, steht zumindest im Standarddeutschen phonetisch

gesehen ein Konsonant davor, nämlich der Glottisverschluss [ʔ], beispielsweise im Wort *iss*: [ʔis] (Maas 2006, S. 124).

Fast alle konsonantischen Phoneme können in betonbaren Silben alleine im Onset stehen. Nicht möglich in betonbaren nativen Silben sind laut Staffeldt (2010, S. 135) die Phoneme [ŋ], [s], [ç], [x]¹. In nicht-nativen Wörtern kann auch manchmal der Laut [ç] auftreten, wie zum Beispiel in *Chemie* (Staffeldt 2010).

Bei einem Onset bestehend aus zwei Konsonanten sind viele verschiedene Varianten möglich, standardmäßig besteht er aber aus einem Obstruenten und einem Sonoranten. Nicht möglich sind zwei hintereinanderfolgende Sonoranten: Der Onset kann genau nur einen Sonoranten enthalten. Nur der Laut [v] stellt eine Ausnahme dar, da dieses Phonem als zwischen den Obstruenten und Sonoranten stehend klassifiziert werden kann. Es kann in erster und zweiter Position im Onset vorkommen, Beispiele hierfür sind die Wörter *Wrack* oder *Schwester* (Duden ⁸2009, S. 43).

In Abbildung 1 sind die Möglichkeiten der Verteilung der Phoneme auf die erste oder zweite Position des Onsets illustriert (Duden ⁸2009, S. 42):

1. Pos. \ 2. Pos.	p	t	k	b	d	g	f	ʃ	v	ts	pf
R	+	+	+	+	+	+	+	+	+		+
l	+		+	+		+	+	+			+
n			+			+		+			
m								+			
y			+					+		+	

Abb. 1.: Möglichkeiten der Verteilung von zwei Konsonanten im Onset (Duden 2009, 8. Auflage, S. 42)

Man erkennt hier, dass sich die stimmlosen Plosive wie ihre stimmhaften Gegenstücke verhalten: [t] verhält sich so wie [d] usw. Bis auf [v] befinden sich, wie gesagt, in der ersten Position nur Obstruenten, in der zweiten nur Sonoranten. Die Affrikaten werden hier monophonematisch gewertet. Bei Fremdwörtern treten Ausnahmen auf:

¹Dem ist hinzuzufügen, dass die Quellen vorwiegend für das Bundesdeutsche sprechen und es im österreichischen Deutsch sehr wohl möglich ist, ein [s] im Onset zu haben, da auch Wörter wie *Sonne* als [ˈsɔnə] ausgesprochen werden können.

Dschungel - [dʒ]. Es handelt sich dabei aber wiederum um eine Obstruent-Sonorant-Abfolge.

Ein Onset mit drei Konsonanten kann nur so aufgebaut sein, dass er mit einem koronalen stimmlosen Frikativ oder einem Affrikaten beginnt und einen stimmlosen Plosiv als zweites Element besitzt. Beispiele dafür sind: [ʃpr], [ʃpl], [ʃtr] (*Sprung, Splint, Strich*) und [tsv], [pfl] (*Zwang, Pflaume*). Ausnahmen hiervon treten in Fremdwörtern auf: [skr], [skl] (*Skrupel, Sklave*) (Staffeldt 2010, S. 135; Duden ⁸2009, S. 42). Es gibt keinen Onset, der mehr als drei Konsonanten beinhaltet. Im Gegensatz dazu schreibt Noack (2010, S. 60), dass im Onset nicht mehr als zwei Phoneme auftreten können, da ein zusätzliches [ʃ] oder [s] ein „extrasilbisches Segment“ ist (vgl. Kapitel 2.2.2 *Die Sonoritätshierarchie* bzw. Kapitel 2.2.3 *Silbengrenzen*).

Hall (2000, S. 244) behauptet, dass die Evidenz für die phonologische Konstituente Onset gering zu sein scheint, da „in den Sprachen der Welt die Zahl der Regeln, die den Onset als Umgebung erfordern, sehr gering“ ist.

2.2.1.2. Nukleus

Der Nukleus kann auch als *Nucleus*, *Kern*, *Gipfel* oder *Silbengipfel* bezeichnet werden. Im Nukleus der Silbe steht im Allgemeinen ein Vokal. Dabei ist zu beachten, dass im Nukleus immer nur ein Segment stehen kann, ein Vollvokal. Bei einem langen Vokal oder einem Diphthong befindet sich das zweite Segment bereits in der Koda der Silbe (Noack 2010). Dies gilt allerdings nur für eine Teilgruppe der Diphthonge, für die **geschlossenen Diphthonge** [aɪ], [ɔɪ], [aʊ], die fest zum Inventar der deutschen Laute gehören. Die andere Gruppe, die **offenen Diphthonge** kommen nur in nicht nativen Wörtern vor (*Region, Union, speziell, sozial, Menuett* usw.) - es handelt sich hier um eine „Folge von nicht silbischem Vokal (z.B. [ʊ] oder [ɪ]) und silbischem Vokal (z.B. [a], [o] oder [ɛ])“ (Duden ⁸2009, S.43). Die nicht silbischen Vokale gelten wiederum nicht als Bestandteil des Silbenkerns, sondern werden oft zum Anfangsrand zugehörend klassifiziert. Sie sind Halbvokale und liegen noch innerhalb der Öffnungsbewegung zum Vollvokal (vgl. Kapitel 2.2.2 *Die Sonoritätshierarchie*). Oft werden sie auch als Konsonanten angesehen und dann als *Gleitlaute* oder *Approximanten* bezeichnet (Duden ⁸2009).

Im Nukleus kann allerdings auch ein Konsonant stehen (ein Sonorant), der dann *silbischer Konsonant* genannt wird und phonologisch folgendermaßen markiert wird:

[l] ist ein silbisches [l], beispielsweise im Wort *segeln*: ['ze:gln].

2.2.1.3. Koda

Die Koda trägt auch die Namen *Coda*, *Endrand*, *Offset*, *Auslaut* oder *Silbenauslaut*.

- Analog zum Onset wird die besetzte Koda *geschlossen* genannt, und sie kann
 - *einfach* (aus einem Laut bestehend) oder
 - *komplex* (aus mehreren Lauten bestehend) besetzt sein;
- oder sie ist unbesetzt und es handelt sich um eine sogenannte *offene* Silbe.

Viele Lautkombinationen des Onsets kommen in der Koda in umgekehrter Reihenfolge vor, zum Beispiel im Onset [kl] und in der Koda [lk] (beispielsweise *klein* und *welk*). Es können höchstens fünf graphematische Konsonanten in der Koda stehen (aber nur vier Phoneme, wenn die Affrikate monophonematisch gewertet werden), zum Beispiel in *schimpfst* ['ʃɪmpfst]. Noack (2010, S. 60) schreibt, dass nur vier Phoneme im Offset auftreten können, da in Beispielen wie *des Herbsts* das [s] als „extrasilbisches Segment“ gezählt wird (vgl. Kapitel 2.2.2 *Die Sonoritätshierarchie*).

Drei wichtige Eigenschaften der Koda sind weiters zu beachten (Duden ⁸2009, S. 44f):

1. Es gibt keine Beschränkung auf nur einen Sonoranten; Liquide und Nasale bilden eine eigene Kategorie, zum Beispiel im Wort *Qualm*, wobei [ʀ] eine höhere Sonorität besitzt als [l] (es existiert das Wort *Kerl* und nicht **Kelr*).
2. In der Koda gibt es keine stimmhaften Obstruenten. Dadurch entsteht die Auslautverhärtung, die sehr charakteristisch für das Deutsche ist. In manchen Varianten tritt aber stattdessen die Spirantisierung auf, bei der [g] zu [ç] wird, und nicht zu [k] (beispielsweise in *lustig* ['lʊstɪç]).
3. Die dritte und die für diese Arbeit wichtigste Eigenschaft der Koda ist der Umstand, dass es zwischen Koda und Nukleus einen Längenausgleich gibt: Wenn die Koda leer ist, ist der Vokal lang und gespannt (betont), wenn zwei oder mehr Konsonanten in der Koda stehen, ist der Vokal kurz und ungespannt. Es gibt allerdings auch viele Ausnahmen dieser Regel, zum Beispiel *Mond* (wo der Vokal lang ist, obwohl die Koda komplex ist), oder *wenn* (wo ein kurzer

Vokal auftritt, obwohl die Koda einfach ist). „Während der Nukleus nur aus einem Segment besteht, können die Ränder komplexer sein.“ (Maas 2006, S. 116).

Aufgrund dieses Zusammenhangs fasst man den Nukleus und die Koda eines Wortes in Kontrast zum Anfangsrand oft zu einer Einheit zusammen. Diese phonologische Komponente wird als *Reim* oder *Silbenreim* bezeichnet und wird im Folgenden besprochen.

2.2.1.4. Die phonologische Komponente „Reim“

Die Bezeichnung Silbenreim erinnert auch daran, dass bei schulmäßig gereimten Versen die letzten Silben wenigstens in Kern und Endrand übereinstimmen (Hut - Mut). (Duden ⁸2009, S. 45)

Der phonologische Reim ist eine subsilbische Struktur, die zwischen der Silbe und den Komponenten Nukleus und Koda angesiedelt ist und aufgrund der Abhängigkeitsbeziehung zwischen Nukleus und Koda auf gleicher Ebene wie der Onset steht. Diese Abhängigkeitsbeziehung wird auch *Dependenz* genannt (Noack 2010):

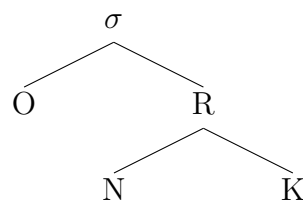


Abb. 2.: Silbenstruktur mit Komponenten σ für Silbe, O for Onset, R für Reim, N für Nukleus und K für Koda (Staffeldt 2010, S. 134)

Die Bezeichnung *Reim* für diese phonologische Einheit, die Nukleus und Koda zusammenfassen soll, kommt aus der Poetik, vom üblichen (End-)Reimverfahren (Maas 2006). Über den Zusammenhang von poetischem Reim und phonologischem Reim schreibt Maas (2006, S.88): „[Der phonologische Reim] beruht auf der Zerlegung der Silbe in ihren sonoren Kern [...] und das komplementäre Element, den Anfangsrand [...] [und liegt dann vor, wenn] Übereinstimmung im sonoren Kern und (eventuell im Endrand) der Silbe bei gleichzeitiger Differenz im Anfangsrand (*Reh* : *Schnee*; *Maid* : *Kleid*) [vorherrscht,] [...] von daher ist es üblich, diese Konstituente als *Reim*

zu bezeichnen.“ Somit ist der Kern des poetischen Endreims die Konstituente der betonten Silbe, die in der silbenstrukturellen Terminologie auch Reim heißt (Maas 2006).

Wiese (²2006) schreibt über die Beziehung zwischen poetischem und phonologischem Reim, dass die beiden durchaus das Gleiche sein können, wenn das letzte Wort einer Zeile in einer betonten Silbe endet. Seiner Meinung nach ist ein zweisilbiger Reim der Poetik (zum Beispiel *Zinnen-Sinnen*) „in keiner Silbentheorie [enthalten]“ (²2006, S. 45, mehr dazu unter Kapitel 2.3 *Mehrsilbige Reime*).

Maas (2006) schreibt, dass die Reimdichtung früher noch toleranter war und es sogar zwei unbetonte Silben nach einer betonten Silbe gab, und führt als Beispiel das Nibelungenlied an, in dem sich die Wörter *Hagene* und *degene* (*Hagen-Helden*) reimten. Wiese (²2006) fügt dem hinzu, dass dies im heutigen Deutsch noch so ist (*heitere-weitere*).

Allerdings ist die Literatur bezüglich der phonologischen Reim-Komponente alles andere als einheitlich. Viele Autoren sprechen sich für oder gegen die Existenz eines phonologischen Reimes aus. Beispielsweise Hall (1992) meint, dass die Argumente für den Reim relativ schwach sind: Das Hauptargument für den Reim ist, dass sogenannte *cooccurrence restrictions* (gemeinsame Restriktionen) zusammen zwischen Nukleus und Koda auftreten, nicht aber zwischen Onset und Nukleus, oder Onset und Koda. Ein Beispiel hierfür wäre im Englischen die Restriktion, dass nach [aʊ] nur koronale Konsonanten kommen (Hall 2000). Aber Hall zeigt, dass es auch Restriktionen gibt, die nicht nur zwischen Nukleus und Koda auftreten können: Er zieht beispielsweise Twaddell (1939, in: Hall 1992) heran, der betont, dass im Englischen eine Abfolge von [pf]-Vokal-velarer Verschlusslaut nicht möglich ist. Twaddell (1940, in: Hall 1992) schreibt auch, dass es eine Tendenz gibt, dass der gleiche Konsonant nicht vor und nach dem Vokal auftritt, und je sonorer der Vokal ist, umso mehr wird dies eingehalten (Hall 1992). Ein Beispiel hierfür wäre das Wort *floor*, im Gegensatz zu **flood* (Hall 2000, S. 241). Beide Male handelt es sich hierbei um eine andere Restriktionsumgebung als um den Reim. Auch der Zusammenhang mit dem poetischen Reim liefert laut Wiese (²2006) keine eindeutige Erklärung: „surprisingly, the poetic rhyme provides rather clear evidence for the onset and for the foot, but not really for the phonological rhyme.“ Die Beispiele, die die Existenz der Komponente Reim belegen wollen, beschränken sich häufig auf einsilbige Wörter. Wenn man mehrsilbige Wörter benutzt, merkt man oft, dass es sich lediglich um den „Post-Onset“-Bereich

2. Was sind Reime?

oder um eine Reihe von Segmenten handelt, die nicht zum phonologischen Reim gehören (Reim + σ). Beispiele hierfür sind: *Techtel-Mechtel*, *higgeldy-piggeldy*, *Hokuspokus*, *heft lemisphere* (Hall 1992, S. 45; Hall 2000, S. 241). Dieses Thema soll aber noch in Kapitel 2.3 *Mehrsilbige Reime* weiter vertieft werden.

Für den Reim spricht laut Hall (1992) die Tatsache, dass es eine starke Reimtradition in der Dichtung, Sprechfehler, Wortspiele usw. gibt, die auf dem Reim basieren könnten. Auch Dell (1986) weist darauf hin, dass Versprecher weit häufiger Nukleus und Koda umfassen, als Onset und Nukleus, „due to the presence of the rime constituent node“ (Dell 1986, S. 312). Wiese (²2006) ist der Meinung, dass die Längenbegrenzung des Vokals genügend Evidenz für den Reim ist. Auch ist der Reim nicht so offensichtlich an phonotaktische Begrenzungen gebunden, gleichzeitig gibt es einige Beschränkungen, die sich auf Nukleus UND Koda beziehen und die Reim-Konstituente betreffen könnten. Beispielsweise wird der Laut /r/, wenn er im Reim vorkommt, und nur dann, vokalisiert - dabei ist es gleichgültig, ob der Laut im Nukleus oder in der Koda vorkommt (Wiese ²2006). Laut Vater (1992) gelten als weiteres Indiz für die Ansetzung einer Reim-Konstituente die Folgen homorganer Nasale, zum Beispiel: *Riemen* [ri:mm̩], *denen* [de:nn̩], *singen* [ziŋŋ̩]. Ein drittes Argument könnte die Auslautverhärtung sein: Sie tritt nicht nur im letzten Konsonanten der Silbe auf, sondern bei Vorhandensein mehrerer Konsonanten hintereinander werden alle Laute stimmlos (Hall 2000). Zusätzlich könnte der Reim als Indiz für das Silbengewicht gelten. Schwere Silben (Silben mit einem langen Vokal oder einem kurzen Vokal und einem Konsonanten) haben einen verzweigenden Reim, leichte Silben (Silben mit einem kurzen Vokal ohne Konsonanten) einen nicht verzweigenden (Hall 2000).

Zusammenfassend kann man sagen, dass Argumente für und gegen die Konstituenten des Reims sprechen. Dies weiter zu erläutern ist allerdings nicht Thema dieser Arbeit. Wichtig hierfür ist zu wissen, wie eine Silbe aufgebaut ist und wo im Wort der Reim genau stattfindet.

Im weiteren wird die Silbe, diese einzelnen Bestandteile zusammenfassend, als Ganzes betrachtet und näher erläutert.

2.2.2. Die Sonoritätshierarchie

“Eine Silbe ist eine Folge von Lauten (im Grenzfall die Einerfolge). Die Abfolge der Laute ist streng geregelt.“ (Duden ⁸2009, S. 38)

Es gibt Beschränkungen dafür, wie eine Silbe aufgebaut ist bzw. welche Laute nacheinander in welcher Reihenfolge auftreten dürfen. Dies basiert auf der sogenannten **Sonoritätshierarchie**, die eng mit der *Phonotaktik* verbunden ist. Die Phonotaktik ist eine „phonologische Teildisziplin, die die Regularitäten der Phonemverbindungen einer Sprache zu Silben, Morphemen und Wörtern zum Gegenstand hat“ (Glück ⁴2010, Stichwort *Phonotaktik*, S. 511). Somit beschreibt sie die „Wohlgeformtheit“ deutscher Lautverbindungen. Die Sonoritätshierarchie zeigt anhand der Sonorität (oder *Schallfülle*) der Laute, welche Laute nacheinander in der Silbe auftreten können, um eine wohlgeformte Silbe zu erzeugen (dabei muss eine theoretisch wohlgeformte Lautabfolge aber keinesfalls auch zwingend existieren).

Die Sonorität nimmt im Onset zu, erreicht im Nukleus ihr Maximum und nimmt zur Koda hin wieder ab (vgl. Kapitel 2.2.1 *Bestandteile der Silbe*), geht also von Lauten mit hohem Geräuschanteil über Laute mit einem hohen Stimmanteil und wieder zurück. Pro Silbe findet dabei im Grunde ein Öffnungs- und ein Schließprozess der Artikulationsorgane statt, wobei beim Nukleus der höchste Grad der Öffnung erreicht ist (Duden ⁸2009). Die Sonoritätshierarchie ist somit eine Skala, die anzeigt, von welchen Lautklassen (Sonoritätsklassen) hier die Rede ist. Obwohl dieses Schema im Grunde für alle Sprachen der Erde gilt und die Sonorität ein universales Prinzip ist, kann die exakte Form des Onsets oder der Koda sprachspezifisch sehr unterschiedlich ausfallen (Noack 2010, Hall 2000). Es werden mindestens fünf Sonoritätsklassen unterschieden (Duden ⁸2009, S. 40):

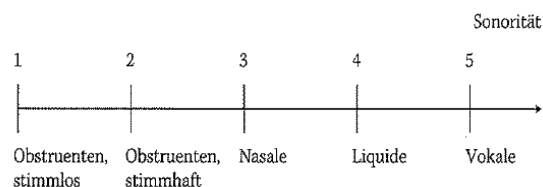


Abb. 3.: Die Sonoritätshierarchie (Duden 2009, 8. Auflage, S. 40)

2. Was sind Reime?

Die maximale Sonorität besitzen die Vokale, die geringste die stimmlosen Plosive. Das Deutsche ist bezüglich der Silbenstruktur und des Aufbaus der Silbenränder ziemlich komplex (Maas 2006). Andere, komplexere Sonoritätshierarchien sind zum Beispiel die nach Jespersen (1904, in: Noack 2010, S. 60):

tiefe	hohe			sth.	sth.	sth.	sth.
Vokale	Vokale	Liquide	Nasale	Frikative	Plosive	Frikative	Plosive
7	6	5	4	3	2	1	0

← zunehmende Schallfülle

Abb. 4.: Die Sonoritätshierarchie nach Jespersen (1904, in: Noack 2010, S. 60)

... beziehungsweise die Sonoritätsskala nach Maas (2006, S. 121):

Rangstufe	Phonetische Qualität	Laute (in IPA)	
1	Vokale, nasal	ã ẽ õ	Vokale
2	Vokale, oral, offen	a ɑ	
3	Vokale, oral, halb-offen	e ε o ɔ ø œ	
4	Vokale, oral, geschlossen	i y u	
5	Halbvokale	ɿ ʏ ʊ	(konsonantische) Sonanten
6	Liquiden	l r ʀ	
7	Nasale	m n ŋ	
8	Frikative, oral, stimmhaft	v ð z ʒ ʒ	Konsonanten (konsonantische Geräuschlaute)
9	Frikative, oral, stimmlos	f θ s ʃ ç x ʁ	
10	Plosive, oral, stimmhaft	b d g	
11	Plosive, oral, stimmlos	p t k	
12	Glottale	h ʔ	

⁴ Nasalvokale (Stufe 1) werden im Deutschen phonologisch nicht genutzt. Halbvokale (Stufe 5) sind im Deutschen relativ offen [ɿ ʏ ʊ], s. Kap. 7; diese Differenzierung ist bei den Vokalen der Stufe 4 (die ungespannten Engevokale [i y u]) nicht aufgenommen, s.

Abb. 5.: Die Sonoritätshierarchie nach Maas (2006, S. 121)

Allerdings gibt es auch ein paar Elemente, die aus dieser Sonoritätshierarchie herausfallen (Hall 2000). Diese werden als *extrasilbisch* bezeichnet. Im Deutschen sind dies die Laute [ʃ] und [s], die eine hohe artikulatorische Intensität besitzen (Noack 2010), wenn sie im Onset vor einem Plosiv oder in der Koda nach einem Plosiv auftreten. Diese extrasilbischen Segmente verhalten sich aber wie Silbenkonstituenten

und beeinflussen die Artikulation. Dieser Umstand ist noch nicht ganz geklärt (Noack 2010).

Maas (2006, S. 124ff) unterscheidet weiters 3 Silbentypen:

- die *prominente Silbe*, die betont ist,
- die *reduzierte Silbe* oder *Reduktionssilbe*, die nicht betont ist und einen reduzierten Nukleus besitzt,
- und die *Normalsilbe/neutrale Silbe* (Noack 2010, S. 56) oder *nicht-prominente, nicht-reduzierte Silbe*, die nicht betont ist, aber einen Vollvokal im Nukleus hat.

Hierbei stellt der dritte Typus eine Ausnahme dar. Im Wort *Tomate* ist die erste Silbe eine solche neutrale Silbe (die zweite ist die prominente Silbe, die letzte die reduzierte Silbe mit dem Schwa). In neutralen Silben kann der Vokal gespannt oder ungespannt sein (wie in den prominenten Silben), aber niemals lang.

2.2.3. Silbengrenzen

„Silbengrenzen treten in Wörtern auf, die aus mehreren Silben bestehen, und ergeben sich aus der Struktur der benachbarten Silben.“ (Duden⁸2009, S. 46)

Es ist nicht immer ganz einfach, die Silbengrenzen festzustellen. Sie können an den gleichen Stellen wie die Morphemgrenzen auftreten, dies ist aber nicht der Regelfall (Staffeldt 2010). Wenn die Silben- und die Morphemgrenze zusammenfallen, handelt es sich um eine morphologisch bedingte Silbengrenze. Wenn dem nicht so ist, dann ist es eine phonologisch bedingte Silbengrenze (Duden⁸2009, Staffeldt 2010). Ein Beispiel hierfür wäre das Wort *Kinder*. Die morphologische Grenze teilt das Wort folgendermaßen auf: *Kind-er*. Anhand der phonologischen Grenze wird es so aufgeteilt: *Kin-der*.

Es können vier Regeln für den Ort der Silbengrenze aufgestellt werden (Duden⁸2009, S. 46). Die Silbengrenze wird im Folgenden in der phonetischen Umschrift mit einem Punkt [.] markiert:

Regel 1 Falls zwei silbische Vokale aufeinander folgen, liegt die Silbengrenze zwischen ihnen (*Ruhe* ['ru:.ə]). Dies gilt nicht für Diphthonge, da es bei Diphthongen ja immer nur ein Element gibt, das das Merkmal [+silbisch] besitzt.

Regel 2 Wenn zwischen zwei Silbenkernen (zwei Vokale oder ein Vokal und ein Diphthong) ein Konsonant ist, gehört dieser Konsonant bereits zur nächsten Silbe und bildet dort den Onset (*Bote* ['bo:tə]).

Regel 3 Betonte Silben mit einem ungespannten Vokal müssen geschlossen sein. Dabei kann ein Konsonant Teil beider Silben werden, zum Beispiel im Wort *Scholle* ['ʃɔlə], das Wort wird unterteilt in [ʃɔ] und [lə]. Diese Konsonanten werden *Silbengelenke* genannt und sind somit ambisilbisch (Altmann und Ziegenhain 2002). Hier ist die Auslautverhärtung nicht wirksam: Wenn ein Plosiv Silbengelenk ist, kann er stimmhaft bleiben (Beispiel *Paddel* ['padəl]). Im Deutschen sind ambisilbische Konsonanten weit verbreitet: Alle Konsonanten außer /h/ und /ʔ/ können ambisilbisch sein (Wiese ²2006, S. 35f).

Regel 4 Wenn zwischen zwei Silben mehrere Konsonanten stehen, werden sie nach Regel 3 aufgeteilt. Falls es mehrere Möglichkeiten geben sollte, dann “gehören alle die Konsonanten zur zweiten Silbe, die zusammen einen wohlgeformten Anfangsrand bilden können“ (Duden ⁸2009, S. 47). Dies ist das Gesetz des maximalen Anfangsrandes (maximal onset principle). Hall (2000, S. 218) meint hierzu: „Bilde zunächst den größtmöglichen Silbenanlaut, dann bilde den Silbenauslaut“. Dieser Vorgang ist aber nicht immer ganz eindeutig. Oft werden Wörter von den Sprechern uneinheitlich syllabiert und manchmal kommen Onsets zustande, die sonst nicht existieren würden, hier zu sehen am Beispiel *Adler*: Es existieren die Varianten ['ʔa:t.lə] und ['ʔa:dlə] (Duden ⁸2009, S. 47).

2.3. Mehrsilbige Reime

In den Studien zu Reimen wird häufig mit einsilbigem Material gearbeitet (DeBleser et al. 2004, Kirtley et al. 1989, Rugg und Barrett 1987, Spencer et al. 2000 uvm.). Wenige Beispiele gibt es für Untersuchungen mit zweisilbigen Reimen (Acheson und MacDonald 2011, Duncan et al. 2007, Inkelas 2003). Aus diesem Grund möchte ich hier näher auf die Reimstruktur in mehrsilbigen Wörtern eingehen.

Jeder Sprecher des Deutschen kann seinem Gefühl nach bestätigen, dass es mehrsilbige Reime gibt, doch gibt es auch eine phonologische Theorie dahinter? Wiese schreibt, dass ein zweisilbiger Reim seines Wissens nach in keiner Silbentheorie enthalten ist (²2006, S. 45). Treiman (1992, in: Duncan et al. 2007) beschreibt ein

mehrsilbiges Wort folgendermaßen:

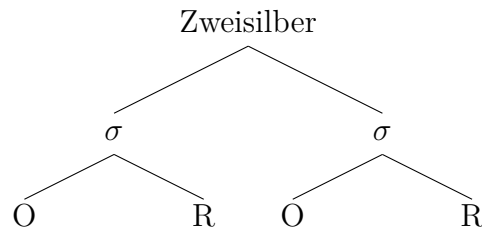


Abb. 6.: Silbenstruktur mit Komponenten σ für Silbe, O for Onset, R für Reim (Treiman 1992, in: Duncan et al. 2007, S. 201)

Nach diesem Modell wird ein zweisilbiges Wort einfach in zwei einzelne Silben unterteilt, die keine weiteren Abhängigkeitsbeziehungen voneinander besitzen. Mehrsilbige Reime allerdings postulieren eher eine Trennung der Silben in Onset und Rest, was auch durch Versprecher unterstützt wird (Davis 1989, in: Duncan et al. 2007). Beispiele, die die Existenz der Reimkomponente in mehrsilbigen Wörtern belegen wollen, können auf eine Wiederholung des Bereichs nach dem Onset beschränkt werden bzw. auf einen Bereich, der nicht die phonologische Komponente Reim in der Silbe umfasst. Ein Beispiel hierfür ist das Wort *Hokuspokus* (vgl. Kapitel 2.2.1.4 *Die phonologische Komponente „Reim“*).

Berg (1992, S.80) allerdings stellt eine Silbenstruktur vor, die zur Beschreibung mehrsilbiger Wörter herangezogen werden kann und den Reim in einzelnen Silben und mehrsilbigen Wörtern erklären kann. Er führt die Komponente „Superreim“ an, hier im Beispiel *Mantel*:

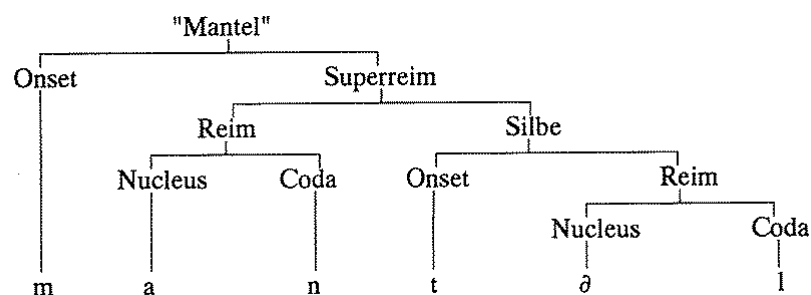


Abb. 7.: Die hierarchische Organisation von Silben in mehrsilbigen Wörtern (Berg 1992, S. 80)

Der „Superreim“ umfasst tatsächlich den Bereich nach dem Onset, erstreckt sich

2. Was sind Reime?

über den Reim der ersten Silbe und über die komplette zweite Silbe, ehe diese weiter in Onset und Reim unterteilt wird. Dies ist eine silbenübergreifende Komponente. Der Superreim von *Mantel* würde somit *-antel* lauten, es reimt auf Wörter mit dem gleichen Superreim, beispielsweise auf *Hantel*. Duncan et al. (2007) schreiben, dass diese Struktur kompatibler mit experimentellen Ergebnissen ist (im Gegensatz zu Abb. 6).

Der Anfang des Superreims tritt allerdings nur bei betonten Silben auf, was die Definition von Reimen auch postuliert hat: Der Reim beginnt bei der letzten betonten Silbe bzw. dem letzten betonten Vokal eines Wortes. Duncan et al. (2007, S. 201) schreibt auch, dass „rhyming words of any size share everything after the onset of the stressed syllables“. Als Beispiele hierfür werden *civility-mobility* und *stationary-inflationary* herangezogen. Wenn vor dieser betonten Silbe noch eine Silbe im Wort vorhanden ist, wird sie einfach „angehängt“. Die Silbenstruktur würde folgendermaßen aussehen, beispielsweise für das Wort *Tomate*:

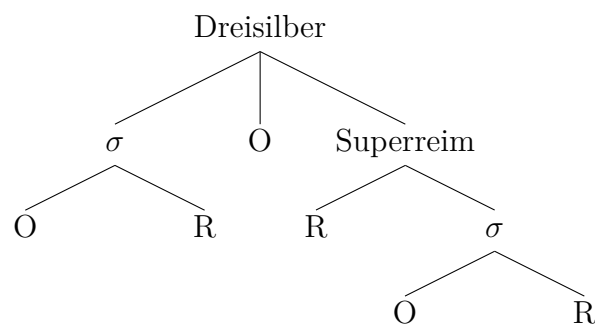


Abb. 8.: Silbenstruktur mit Komponenten σ für Silbe, O for Onset, R für Reim (Berg 1989, in: Duncan et al. 2007, S. 202)

Evidenz für eine Superreimtheorie im Gegensatz zu einer Einzelsilbentheorie liefern verschiedene Studien (vgl. Kapitel 3.1 *Silben und Reime im Spracherwerb* bzw. Kapitel 3.2 *Die phonologische Bewusstheit*).

3. Überblick über den aktuellen Forschungsstand zu Reimen in der Psycho-, Patho- und Neurolinguistik

Dieses Kapitel versucht einen Einblick in das breite Forschungsgebiet der Reime in der Psycho-, Patho- und Neurolinguistik zu geben. Zunächst soll vom Kleinkindalter an die Rolle der Reime und Silben im Spracherwerb beleuchtet werden, um dann auf den Schriftspracherwerb und die phonologische Bewusstheit einzugehen. Diese Forschungsgebiete werden nur angerissen, da sie äußerst umfangreich sind und im Rahmen dieser Arbeit nicht näher bearbeitet werden können. In jedem Unterkapitel befinden sich viele Verweise auf weiterführende Literatur. Es folgt die Verarbeitung von Reimen bei sprachgesunden und sprachbeeinträchtigten Erwachsenen, bis schließlich auf die neurologische Verarbeitung von Reimen eingegangen wird.

3.1. Silben und Reime im Spracherwerb

In Kapitel 2 *Was sind Reime?* haben wir bereits gesehen, wie eng Reime und Silben zusammenhängen. Staffeldt (2010) schreibt, dass jeder kompetente Sprecher einer Sprache über die Fähigkeit verfügt, Silben identifizieren zu können. Dies gilt für Erwachsene, sowohl auch wie für Kinder:

„Der Erwerb der Fähigkeit, silbisch sprechen zu können, ist der erste wirklich große Schritt beim Erwerb der Muttersprache.“ (Staffeldt 2010, S. 133)

Die ersten fünf Monate im Leben eines Babys nennt man *Vorsilbenalter*, danach kommt der Zeitraum des *Silbenplapperns* im Zeitraum von sechs bis zwölf Monaten

(Staffeldt 2010). Kinder vor der Einschulung werden auf die Frage, aus welchen Einheiten ein mehrsilbiges Wort besteht, mit den Silben antworten, während sie noch nicht wissen, dass Wörter oder eben jene Silben aus mehreren einzelnen Lauten aufgebaut sind (Staffeldt 2010).

„[Bei Silben] handelt es sich um die kleinste intuitiv wahrnehmbare Sprach-einheit, die unabhängig von gesteuertem Sprachunterricht kognitiv verfügbar ist. Demgegenüber werden Einzellaute - das haben zahlreiche Experimente mit Schreibanfängern (z.B. Bosch 1937, Röber 2009) gezeigt - vor dem Schrifterwerb i.d.R. noch nicht wahrgenommen.“ (Noack 2010, S. 50)

Dies erklärt Noack (2010) dadurch, dass eine sprachliche Äußerung nicht einfach aus unterschiedlichen, voneinander abgetrennten Lauten besteht, sondern aus einem Schallkontinuum, das ständig fließende Übergänge beinhaltet (Öffnungs- und Schließbewegung der Artikulationsorgane).

Auch in Liedern und Gedichten werden die Silben als rhythmische Einheiten benutzt, da auf ihnen der Akzent liegen kann. Die Unterteilung des Sprachflusses in Silben ist für Kinder und analphabetische Erwachsene einfach und natürlich, während eine andere Unterteilung, beispielsweise in Laute oder gar Morpheme, sich als sehr schwierig gestaltet (Wiese ²2006).

Silben spielen in Sprachspielen „eine zentrale Rolle“ (Berg 1992, S. 69). Wenn die Kinder Fehllassoziationen herstellen, sogenannte *Malapropismen*, dann gibt es einen engen formalen Zusammenhang zwischen dem richtigen und dem nachgemachten, falschen Wort oder der falschen Phonemfolge. Zum Beispiel sagt das Kind „Trampolin“ statt „Tambourin“ - diese beiden Wörter stimmen in der Silbenanzahl überein (Berg 1992, S. 70f). Solche Fehllassoziationen, die bezüglich ihrer Silbenanzahl mit dem gesuchten Wort übereinstimmen, kommen bei Kindern zu 84 % vor, bei Erwachsenen nur bis zu 67 % (Aitchison und Straf 1981, in: Berg 1992). Kinder neigen also eher dazu, die Silbenanzahl beizubehalten. Nach Berg (1992) besitzen sie eine besondere Silbensensitivität.

Duncan et al (2007, S. 200) postulieren, dass auch der Reim ab einem frühen Alter prominent ist, was durch „Kinderreime [...], die spontane Reimproduktion [der Kinder] im Spracherwerb [...] und ihre Fähigkeit, auf Wunsch reimende Gedichte

zu entwerfen [...] [bestärkt wird]“ und geben weitere Quellen für eine frühe Reimbewusstheit an. Die Reimbewusstheit (Unterteilung der Silbe in Onset und Reim, siehe Kapitel 2.2 *Die Silbe*) entwickelt sich womöglich nach der Silbenbewusstheit (Treiman und Zukowski 1996, in: Stadie et al. 2003).

Kirtley et al. (1989) schreiben, dass Kinder vor der Einschulung bereits Reimbewusstheit besitzen: Sie können Silben ganz natürlich in Onset und Reim unterteilen, finden es aber schwierig, mit einzelnen Phonemen zu arbeiten (Stadie et al. 2003). Die Autoren fanden heraus, dass Kinder vor der Einschulung einzelne Phoneme besser wahrnehmen konnten, wenn sie mit dem Onset einer Silbe (bei einsilbigen Wörtern) übereinstimmten, als wenn sie mit einem Teil des Reims übereinstimmten. Die Reimfähigkeit verbessert sich noch vor der Einschulung und ist bei Produktionsaufgaben der Alliteration überlegen (Duncan und Seymour 2000; Seymour und Evans 1994, beide in: Duncan et al. 2007).

Inkelas (2003) hat die Reimproduktion bei einem Kleinkind untersucht, und ist dabei auch auf zweisilbige Reime eingegangen. Im Alter von zwei Jahren basierte ein zweisilbiger Reim auf dem Superreim, wenn der Hauptakzent auf der ersten Silbe lag und der alleinige Akzent des Wortes war. Fiel allerdings der Haupt- oder Nebenakzent auf die zweite Silbe (*plácemât*), so wurde nur die Reimkomponente der zweiten Silbe gereimt (*bat*). Im Alter von vier Jahren wurde die zweite Silbe nur dann alleine gereimt, wenn sie den Hauptakzent bekam, beim Nebenakzent nicht mehr (*sídewàlk* - *bidewalk*), was dem Reimverhalten Erwachsener entspricht. Die Komponente des Superreims gewinnt an Bedeutung. Inkelas erklärt dieses Entwicklungsmuster durch die Reime, die in Geschichten vorkommen, die das Kind hört, und durch den steigenden Einfluss der linguistischen Komponente Reim, oder durch die Reifung der phonologischen Repräsentationen von mehrsilbigen Wörtern.

Nun haben wir die Silben- und die Reimbewusstheit bei Kindern vor der Einschulung beleuchtet. Im Folgenden wird der Begriff der phonologischen Bewusstheit behandelt, der mit dem Schriftspracherwerb in Zusammenhang steht, und es wird besprochen, inwiefern sich die Reimbewusstheit weiter entwickelt.

3.2. Die phonologische Bewusstheit

Die Fähigkeit, Reime zu erkennen und zu bilden ist ein Teilbereich der *phonologischen Bewusstheit*. Wagensveld et al. (2013) schreiben, dass die Reimbewusstheit

eine der ersten Formen der phonologischen Bewusstheit darstellt. Die phonologische Bewusstheit beschreibt die Fähigkeit vom Verständnis des Aufbaus der Wörter. Ptok et al. (2008) fassen den Begriff der phonologischen Bewusstheit wie folgt zusammen:

„Unter diesem Begriff werden die Fähigkeit verstanden, die Aufmerksamkeit unabhängig vom Inhalt bzw. der Bedeutung des Gesagten auf die formal-lautlichen Aspekte der Sprache zu richten, sowie das Bewusstsein darüber, dass man gesprochene Sprache in kleinere Einheiten zerlegen, mit diesen operieren und sie auch wieder zu größeren, komplexeren Einheiten verbinden kann.“ (Ptok et al. 2008, S. 860)

Die phonologische Bewusstheit gilt als Indikator für einen „erfolgreichen Einstieg in das Lesen- und Schreibenlernen“ (Noack 2010, S. 92; siehe auch Stadie et al. 2003 für weitere Literaturhinweise). Andere Autoren meinen, dass erst der Erwerb der Schriftsprache die Entwicklung der phonologischen Bewusstheit begünstigt und vorantreibt (Stadie et al. 2003). Die Autoren scheinen sich aber einig darüber zu sein, dass es einen engen Zusammenhang zwischen phonologischer Bewusstheit und Schriftsprache gibt.

Ptok et al. (2008) schreiben, dass sich „mit der Überprüfung der phonologischen Bewusstheit die spätere Lese-Rechtschreib-Kompetenz vorhersagen lässt [...] und dass durch entsprechende Intervention die späteren Schriftsprachkompetenzen positiv beeinflusst werden können“ (Ptok et al. 2008, S. 862). Kinder mit einer Lese-Rechtschreib-Störung zeigen eine schlechtere phonologische Bewusstheit als unbeeinträchtigte gleichaltrige Kinder (Klicpera und Gasteiger-Klicpera 1994, in: Ptok et al. 2008).

Aufgaben bzw. Fähigkeiten zur phonologischen Bewusstheit könnten beispielsweise umfassen (Stadie et al. 2003, S. 146ff.):

- Auditives Entscheiden der Wortlänge
- Detektieren von Silben oder Phonemen in Wörtern
- Detektieren von Reimpaaren in Wörtern oder Nichtwörtern
- Reimurteil nach Bildvorgabe
- Mündliches Produzieren von Reimen

- Zusammenfügen von Phonemen oder Silben
- Nachsprechen rückwärts von Silben oder mehreren Phonemen (zum Beispiel bei einem Nichtwort wie [y:l])

Weber-Fox et al. (2003) fassen die Reimfähigkeiten im Rahmen der phonologischen Bewusstheit und im Zusammenhang mit dem Schriftspracherwerb wie folgt zusammen:

„The ability to judge whether two words rhyme is one of several measures of phonological awareness that include such skills as phoneme identification, syllable counting, and phonological memory [...]. Children typically develop rhyme awareness during the preschool years, as early as 3 years of age [...]. The abilities necessary to make rhyming judgments for pairs of words (i.e., prime-target) include retrieving the phonological representation of a prime, holding it in working memory, and segmenting it into its onset and rime constituents. The same processes would then be utilized for processing the target word in the pair. To complete the rhyme judgment, the rime constituent of the prime must be compared to the rime of the target. This type of phonological awareness has been shown to be one of several factors that predict the development of reading abilities in children [...]. In fact, normal development of rhyming awareness, in addition to phonemic awareness, is a seemingly crucial skill for the subsequent acquisition of good reading skills [...].“ (Weber-Fox et al. 2003, S. 128)

Bryant et al. (1989, in: Goswami und Mead 1992) haben gezeigt, dass es einen Zusammenhang gibt zwischen dem Wissen um Kinderreime im Alter von drei Jahren und der Lesefähigkeit im Alter von sechs Jahren. Somit scheinen Reimkenntnisse auf den Erwerb der Schriftsprache einzuwirken - doch auch die Schriftsprache kann zurück auf die Reimfähigkeit wirken: Duncan et al. (2007) haben untersucht, ob die Einschulung bzw. der Erwerb der Schriftsprache einen Einfluss auf die Reimfähigkeit von Kindern haben würde. Die Ergebnisse zeigten, dass sich die Reimfähigkeit für einsilbige Wörter von der ersten zur zweiten Klasse Grundschule beinahe um die Hälfte verdoppelt hat. Die Antworten der Kinder der zweiten Klasse und der Erwachsenengruppe waren kaum zu unterscheiden, während die Kinder der ersten

Klasse noch mehr Fehler machten. Man kann daraus schließen, dass der Schriftspracherwerb einen großen Einfluss auf die Verarbeitung und Repräsentation von Reimen bei Kindern hat und dass sich die Reimfähigkeit sehr früh entwickelt und auch bald abgeschlossen ist.

Doch wie weit reicht die Repräsentation eines Reimes bei einem Sprecher? Duncan et al. (2007) untersuchten die Reimproduktion von zweisilbigen Wörtern mit unterschiedlichem Betonungsmuster bei Kindern (erste bis dritte Klasse) und Erwachsenen. Den Probanden wurden einzelne Wörter präsentiert und sie sollten Reimwörter dazu bilden. Die Hypothese der Autoren war, dass die Reimproduktion von zweisilbigen Wörtern die interne Struktur dieser Wörter wiedergeben würde. Würde die Betonung auf der ersten Silbe liegen, würde bei der Reimproduktionsaufgabe der Superreim (vgl. Abb. 7 in Kapitel 2.3 *Mehrsilbige Reime*) wiederholt werden, läge die Betonung auf der zweiten (und letzten) Silbe, würde nur der Reim dieser letzten Silbe auftreten. Zusätzlich meinten die Autoren, dass die Wahl der Reimantwort (Superreim, Silbe oder Reim) auch von der Anzahl der phonologischen Nachbarn abhängen würde, die einen Superreim, eine Silbe oder einen Reim mit dem Zielwort teilen. Das Experiment wurde mündlich durchgeführt. Sie fanden heraus, dass die häufigste Antwort beider Probandengruppen bei Wörtern mit Betonung auf der ersten Silbe eine Superreim-Antwort war. Zweisilber mit Betonung auf der letzten Silbe zeigten ein ungleiches Muster, die Antworten waren nicht eindeutig zu klassifizieren: Superreim-Antworten waren dabei geringer, traten aber immer noch auf, Antworten, die sich mit der phonologischen Komponente Reim oder der Silbe (also Reim plus Onset) reimten, nahmen zu. Dieses Experiment wiederholten die Autoren mit Nichtwörtern. Hier hat sich gezeigt, dass bei Wörtern mit Betonung auf der ersten Silbe nur Superreime produziert wurden. Bei Wörtern mit Betonung auf der zweiten Silbe gab es ein variables Muster wie im ersten Experiment: Superreime, phonologische Reime und Silben traten als Antwort gleichermaßen auf. Die Ergebnisse der Experimente decken sich für zweisilbige Wörter mit Betonung auf der ersten Silbe mit dem Silbenmodell von Berg (1989, in: Duncan et al. 2007, vgl. dazu Abbildung 7 auf Seite 17). Bei Wörtern mit Betonung auf der zweiten Silbe ist die Definition eines Reims weniger eindeutig. Sie fanden auch heraus, dass es nicht die Dichte der phonologischen Nachbarn war, die die Reimantwort beeinflusste.

3.3. Reime bei sprachgesunden Erwachsenen

Bei Erwachsenen sind der Sprach-, sowie der Schriftspracherwerb bereits abgeschlossen, sie besitzen einen höheren Grad phonologischer Bewusstheit als jüngere Menschen bzw. Kinder vor der Einschulung.

Gordon und Baum (1994) schreiben, dass Reim-Priming (vgl. Kapitel 5.1 *Priming-Experimente*) das lexikalische Entscheiden bei visuellen und auditiven Aufgaben mit Wörtern und Nichtwörtern begünstigen.

Nickels und Cole-Virtue (2004) führten eine Studie mit Kontrollprobanden durch, die die Fähigkeiten von unbeeinträchtigten Sprechern mit den Fähigkeiten von Aphasikern vergleichen sollte. Im englischen Sprachraum gibt es ein Diagnostikverfahren namens PALPA (Psycholinguistic Assessments of Language Processing in Aphasia; Kay et al. 1992, in: Nickels und Cole-Virtue 2004), für das es bis dato keine Untersuchungen zur Leistung von unbeeinträchtigten Sprechern gab. Ein Untertest (Subtest 15) überprüfte das Reimentscheiden zweier schriftlich präsentierter Wörter, bestehend aus 60 Wortpaaren.

„To complete this task correctly the participant has to derive phonology from the written word, segment off the rime, and compare the segmented stimuli.“ (Nickels und Cole-Virtue 2004, S. 105)

Die Stimuli waren folgendermaßen aufgebaut:

- reimende Wortpaare
 - reguläre Schreibweise: Eine richtige Entscheidung konnte rein aufgrund des Schriftbildes getroffen werden (*town* - *gown*).
 - nicht reguläre Schreibweise: Eine richtige Antwort konnte nur gegeben werden, wenn die Phonologie des Wortes bekannt war (*bowl* - *mole*).
- nicht reimende Wortpaare
 - reguläre Schreibweise: Eine Entscheidung aufgrund des Schriftbildes führt zu einer nicht korrekten Antwort (*down* - *flown*).
 - irreguläre Schreibweise: Die Wörter besaßen kein gleiches Schriftbild (*hoe* - *chew*).

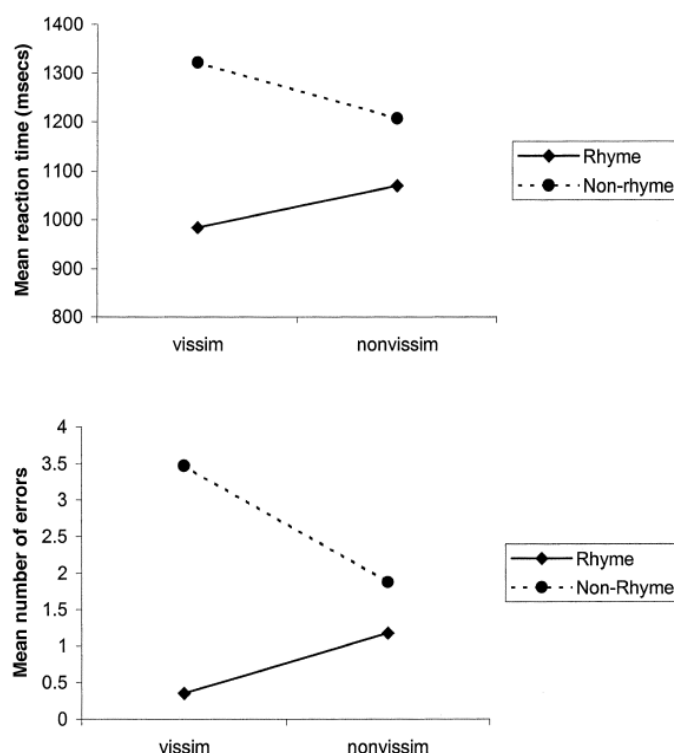


Abb. 9.: Reaktionszeit und Fehlerrate beim Reimentscheiden bei gesunden Erwachsenen vissim = ähnliches Schriftbild, nonvissim = nicht ähnliches Schriftbild (Nickels und Cole-Virtue 2004, S. 113)

Wie anhand von Abbildung 9 zu erkennen ist, gab es eine signifikante Interaktion zwischen Reimen und Schriftbildern: Die Reaktionszeit und die Fehleranzahl stiegen, wenn die Reimpaare ein nicht gleiches Schriftbild besaßen, oder wenn Nichtreimpaare ein gleiches Schriftbild besaßen. Reime mit ähnlichem Schriftbild wurden am schnellsten verarbeitet, es folgten Reime mit unähnlichem Schriftbild, dann Nichtreime mit unähnlichem Schriftbild, und Nichtreime mit ähnlichem Schriftbild wurden am langsamsten verarbeitet. Abbildung 9 zeigt die Fehlerquote pro Untergruppe der Reimpaare mit jeweils 30 Reimpaaren (ähnliches Schriftbild und nicht ähnliches Schriftbild).

Im Durchschnitt beurteilten die Probanden bei den 60 Reimpaaren somit nur 53,12 Wortpaare korrekt. Man darf also bei der Bewertung von sprachbeeinträchtigten Patienten (wozu das Diagnostikmaterial PALPA erstellt wurde) nicht ohne Überprüfung davon ausgehen, dass unbeeinträchtigte Personen die Aufgaben immer korrekt

lösen würden. Nickels und Cole-Virtue schreiben, dass die Probanden bei der Reimentscheidung eine „schwache Leistung“ zeigten (Nickels und Cole-Virtue 2004, S. 114).

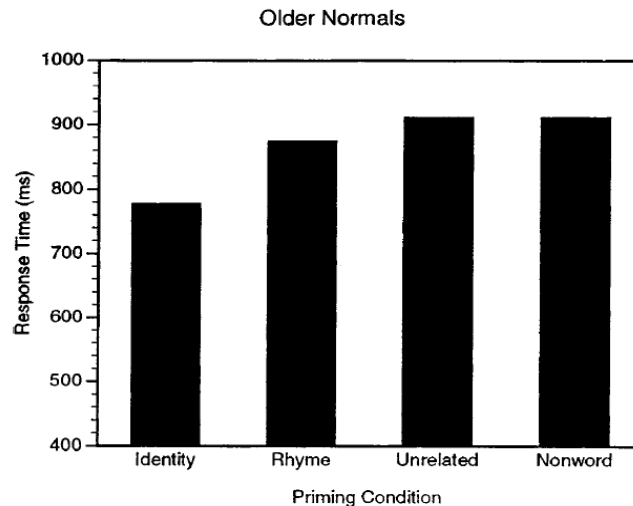


Abb. 10.: Reaktionszeiten beim lexikalischen Entscheiden bei unbeeinträchtigten Probanden in vier verschiedenen Bedingungen (Wortpaare Prime-Zielwort) (Blumstein et al. 2000, S. 83)

Blumstein et al. (2000) untersuchten in einem Priming-Experiment, welche Art von Prime die Verarbeitung des zweiten Stimulus' bei einer lexikalischen Entscheidungsaufgabe verändert. Einmal war der Prime identisch, einmal ein Reimwort, einmal unrelatiert und einmal ein unrelatiertes Nichtwort (Näheres dazu unter 3.4 *Reime bei aphasischen Erwachsenen*). Es hat sich gezeigt, dass der zweite Stimulus in einem identischen Wortpaar als schnellstes beurteilt wurde. Danach folgte die Beurteilung des zweiten Stimulus in einem Reimpaar, die zweiten Stimuli bei den unrelatierten Bedingungen wurden langsamer beurteilt (vgl. dazu Abbildung 10).

Acheson und MacDonald (2011) untersuchten Reime in einem größeren Kontext: innerhalb eines Satzes. Im Gegensatz zur allgemein angenommenen Meinung, dass „die phonologischen Eigenschaften eines Wortes wenig Einfluss auf das Satzverständnis über die Prozesse der Worterkennung hinaus“ (Acheson und MacDonald 2011, S. 193) besäßen, haben die Autoren herausgefunden, dass die Phonologie sehr wohl die Speicherung und Interpretation von Sätzen beeinflussen kann. Die Stimuli bestan-

den aus 24 Satzpaaren mit eingebetteten Objektrelativsätzen. In einem Satz eines jeden Paares bestand eine große phonologische Ähnlichkeit zwischen dem Nomen des Hauptsatzes und dem direkten Objekt des Relativsatzes, sowie zwischen dem Verb des Hauptsatzes und dem Verb des Relativsatzes. Ein Beispiel hierfür wäre:

- Version a: *The baker that sought the banker bought the house.*
- Version b: *The runner that feared the banker bought the house.*

Immer nur das Nomen des Hauptsatzes und das Verb des Relativsatzes wurden in den beiden Bedingungen verändert.

In einem weiteren Experiment, dessen Ergebnisse ich hier zusammengefasst mit denen des ersten Experiments besprechen möchte, wurden die Objektrelativsätze zu Subjektrelativsätzen umgewandelt. Wieder waren nur das erste Nomen sowie das eingebettete Verb die veränderten Variablen in den beiden Bedingungen:

- Version a: *The baker that the banker sought bought the house.*
- Version b: *The runner that the banker feared bought the house.*

Es gab kein Kriterium bezüglich phonologischer Überlappung, manchmal reimten sich die Wörter (wie hier bei *bought - sought*), manchmal besaßen sie nur einen großen Anteil gleicher oder ähnlicher Phoneme (wie in *baker - banker*).

Zu jedem Satz gehörten zwei Verständnisfragen (eine mit einer Bejahung und eine mit einer Verneinung als korrekte Antwort), zum Beispiel: *Did the baker seek the banker?* und *Did the runner buy the house?* bei den b-Versionen, oder *Did the banker buy the house?* bei der a-Version.

Die Probanden sollten die Sätze in ihrer eigenen Geschwindigkeit lesen. Sie mussten zum Erscheinen des nächsten Wortes eine Taste drücken, wodurch ihre Lesegeschwindigkeit gemessen wurde. Nachdem der gesamte Satz gelesen wurde, erschien die Frage, die sie bejahen oder verneinen sollten, und sie bekamen Feedback bezüglich der Richtigkeit ihrer Antworten.

Die Ergebnisse zeigen, dass die Sätze mit großer phonologischer Ähnlichkeit langsamer gelesen wurden. Bei den Objektrelativsätzen gab es besonders langsame Lesezeiten bei den beiden Verben. Somit hat die phonologische Überlappung einen Effekt auf die Lesezeit. Bezüglich den Verständnisfragen wurde herausgefunden, dass die Probanden bei den b-Versionen (ohne phonologische Ähnlichkeit) richtigere und

schnellere Antworten gaben. Phonologische Überlappung in Wörtern innerhalb eines Satzes haben also einen Effekt auf on-line-Messungen (Lesezeit) und off-line-Messungen (Verständnisfragen - bezüglich Richtigkeit und Antwortzeit), und dies sowohl bei Objekt-, als auch bei Subjektrelativsätzen.

Der Effekt phonologischer Ähnlichkeit wurde bereits in Tests zum Kurzzeitgedächtnis genauer untersucht, wo sich auch gezeigt hat, dass Wörter mit großer phonologischer Ähnlichkeit schwieriger abzurufen sind. Das Gedächtnis wurde durch die Reihenfolge, in der die Wörter präsentiert wurden, beeinflusst, nicht durch die Wörter selbst (Acheson und MacDonald 2011).

Aufgrund dieser Ergebnisse stellt sich die Frage, ob Personen nicht vermeiden wollen, dass Reime und andere phonologische Ähnlichkeiten in zwei Wörtern nah zueinander bzw. innerhalb eines Satzes auftreten sollen. Dies würde weitere Konsequenzen für das Experiment haben, das im Rahmen dieser Arbeit durchgeführt wird (vgl. Kapitel 5 *Experiment*).

Einen weiteren Untersuchungsgegenstand bei Erwachsenen im Zusammenhang mit Reimen stellt die Studie von McGlone und Tofighbakhsh (2000, S. 424) dar: „We explored the role that poetic form can play in people’s perceptions of the accuracy of aphorisms.“ Den Probanden wurden Aphorismen präsentiert, die sie bezüglich ihrer Glaubhaftigkeit bewerten sollten. Jeder sich reimende Aphorismus besaß eine sich nicht reimende Kontrollversion (McGlone und Tofighbaksh 2000, S. 426):

- reimender Aphorismus: *Caution and measure will win you treasure.*
- nicht reimender Aphorismus: *Caution and measure will win you riches.*

Den Probanden wurde je eine von zwei Instruktionen gegeben: Bei der ersten Instruktion sollten die Probanden bewerten, inwiefern der (schriftlich) präsentierte Aphorismus ein menschliches Verhalten korrekt beschreibt, bei der anderen Version der Instruktion wurden die Probanden explizit darauf hingewiesen, dass sie ihr Urteil NUR bezüglich der korrekten (semantischen) Beschreibung des Verhaltens fällen sollten, nicht bezüglich der poetischen Qualität der Formulierung. Es hat sich gezeigt, dass sich reimende Aphorismen als korrekter wahrgenommen wurden als sich nicht reimende Aphorismen, aber nur bei der ersten Version der Instruktion. Wenn die Probanden darauf hingewiesen wurden, bei der Bewertung nicht die Formulierung zu integrieren, reduzierte sich dieser Effekt signifikant. Am Ende der Untersuchung

wurden die Probanden gefragt, ob sie glaubten, dass reimende Aphorismen wahrer sind als nicht reimende, was alle Probanden mit „nein“ beantworteten. Somit kann man darauf schließen, dass sich durch (unbewusst wahrgenommene) Reime die Glaubhaftigkeit des Gesagten erhöht, wenn es in einem passenden Kontext präsentiert wird - bei offensichtlicher Unwahrheit wird auch der Reim die Glaubhaftigkeit nicht erhöhen (mehr dazu bei McGlone und Tofighbaksh 2000).

3.4. Reime bei aphasischen Erwachsenen

Die breite Anwendung von Reimen geht sogar über den Bereich der unbeeinträchtigten Sprache hinaus. Hier werden ein paar Studien vorgestellt und Ergebnisse präsentiert, die sich mit den Prozessen bei der Verarbeitung von Reimen bei aphasischen Patienten befassen. Zunächst soll es von der Verarbeitung von Reimen in Priming-Experimenten (mehr über Priming siehe Kapitel 5.1 *Priming-Experimente*) handeln, um später auf die Produktion von Reimen überzugehen.

Blumstein et al. (2000) schreiben, dass der Einfluss phonologischer Faktoren beim Zugriff aufs Lexikon in letzter Zeit im Fokus des Interesses in der Aphasieforschung liegt: „[...] the basis of impairments for aphasic patients may lie in those processes that map sound structure on to lexical form“ (Blumstein et al. 2000, S. 76), und nicht an dem Lexikon oder der Phonologie selbst. Sie führten ein Experiment durch, um zu untersuchen, ob Patienten mit Broca- und Wernicke-Aphasien Reim-Priming so wie unbeeinträchtigte Personen zeigten. In diesem Experiment gab es drei Probandengruppen: 12 unbeeinträchtigte Kontrollprobanden (gematcht im Alter), 9 Broca-Aphasiker und 8 Wernicke-Aphasiker. Die Stimuli haben Blumstein und Kollegen von Burton (1989, in: Blumstein et al. 2000) übernommen und bestanden aus Wörtern und Nichtwörtern: Es gab identische Wortpaare (*take - take*), reimende Wortpaare (*sake-take*), unrelatierte Wortpaare (*bush-take*), unrelatierte Nichtwortpaare (*moob-take*), sowie Füllerpaare (bestehend aus jeweils nur Nichtwörtern und nur Wörtern). Es handelte sich um eine auditive Aufgabe zum lexikalischen Entscheiden, die Probanden und Patienten mussten entscheiden, ob das zweite Wort ein Wort oder ein Nichtwort war. Die Theorie dahinter besagt, dass die Präsentation eines lexikalischen Items nicht nur das Item selbst aktiviert, sondern auch phonologisch ähnliche Items im Lexikon. Somit ist bei Präsentation eines Reimwortes dieses Wort durch seinen vorausgegangenen Prime bereits teilweise aktiviert, wodurch sich die Zeit des Lexika-

litätsurteils verringert und Priming bzw. ein Priming-Effekt stattfindet. Aus diesem Grund geschieht Identity Priming noch schneller.

Die Ergebnisse der Aphasiker sind in Abbildung 11 (im Vergleich zu Abbildung 10 für unbeeinträchtigte Personen auf Seite 27) abgebildet:

- Bei den unbeeinträchtigten Probanden waren die Reaktionszeiten bei den identischen Wortpaaren und den Reimpaaren signifikant schneller als bei den unrelatierten Wortpaaren. Der Reimpriming-Effekt war signifikant langsamer als der Identity Priming-Effekt.
- Bei den Broca-Aphasikern zeigte sich eine signifikant langsamere Reaktionszeit bei Reimen im Vergleich zu identischen und unrelatierten Wortpaaren. Es gab keinen signifikanten Unterschied zwischen identischen und unrelatierten Wortpaaren, daher zeigte sich kein Identity Priming-Effekt. Bei den Reimpaaren wurde eine Hemmung festgestellt, es gab also keinen phonologischen Priming-Effekt.
- Bei den Wernicke-Aphasikern waren die Reaktionszeiten generell langsamer als bei den Broca-Aphasikern in allen Bedingungen. Es gab einen signifikanten Identity Priming-Effekt und einen Reimpriming-Effekt im Vergleich zur unrelatierten Bedingung, und der Identity Priming-Effekt war signifikant schneller als die Reimbedingung, somit verhalten sich die Wernicke-Aphasiker hier ähnlich wie die unbeeinträchtigten Personen (vergleiche dazu Abbildung 10 auf Seite 27).

Broca-Aphasiker machen mehr Fehler als Wernicke-Aphasiker bei den Reimpaaren und den identischen Wortpaaren, somit in den Bedingungen, in denen die Phonologie eine wichtige Rolle spielt.

„The deficits for Broca’s aphasics appear to lie in the mapping of phonological form onto a lexical representation and, in particular, on those processes influencing the activation of lexical candidates from phonological form“ (Blumstein et al. 2000, S. 86).

Eine Erklärung für das negative Priming bzw. die Hemmung bei Reimpaaren könnte in einer reduzierten Aktivierung der lexikalischen Kandidaten liegen. Sie zeigen keine verkürzte Zeit beim lexikalischen Entscheiden, da ein phonologisch ähnliches Wort zu

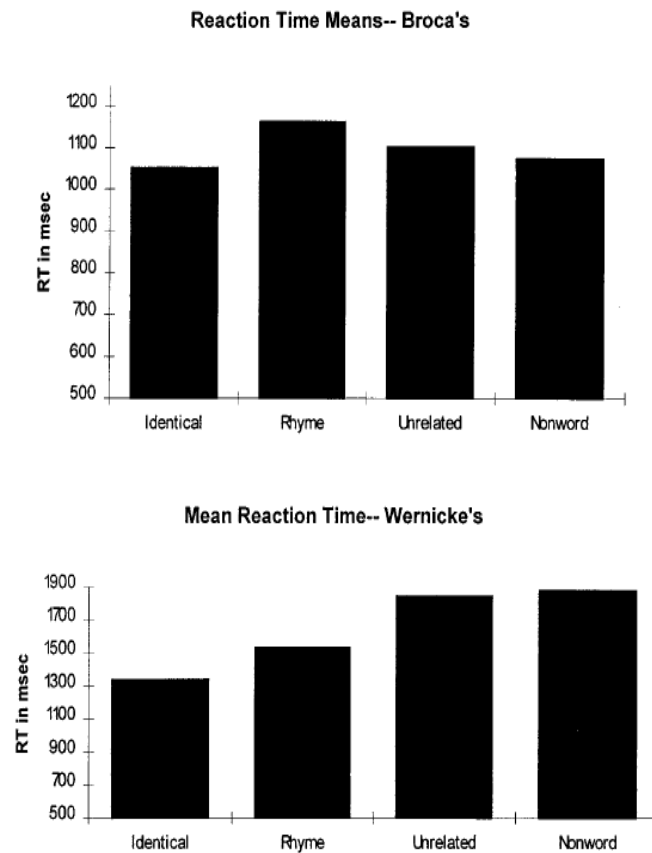


Abb. 11.: Reaktionszeiten beim lexikalischen Entscheiden bei Broca- und Wernicke-Asphasikern in vier verschiedenen Bedingungen (Wortpaare Prime-Zielwort) (Blumstein et al. 2000, S. 84)

einem Prime nicht aktiviert wird. Der Vorgang beim negativen Reim-Priming könnte auf zwei Arten geschehen: Entweder wird ein bestimmtes Wort nicht genügend aktiviert, die phonologischen Nachbarn werden auch aktiviert und müssen dann wieder gehemmt werden. Die Verzögerung entsteht dadurch, dass das richtige Wort aus den konkurrierenden phonologischen Nachbarn herausgefiltert werden muss. Oder das Zielwort sowie die Nachbarn werden nicht genügend aktiviert, wodurch es länger dauert, bis sie die Aktivierungsschwelle übertreten.

Bei den Wernicke-Aphasikern liegt das Defizit nicht so sehr beim Mapping der phonologischen Form auf lexikalische Repräsentationen, sondern bei der Aktivierung der lexikalischen Repräsentation und die Verteilung der Aktivierung im lexikalischen

Netzwerk.

Im Kontrast dazu haben Gordon und Baum (1994) herausgefunden, dass Patienten mit der Diagnose einer nicht flüssigen Aphasie Reimpriming zeigten, während Patienten mit flüssiger Aphasie kein Reimpriming zeigten.

In der Studie von Gordon und Baum (1994) werden Reime beim lexikalischen Verarbeiten bei normalen Probanden (12 Personen) untersucht, sowie bei aphasischen Patienten mit Unterteilung in flüssige Aphasie (ausgezeichnet durch semantische Paraphasien, 9 Patienten) und nicht flüssige Aphasie (mit phonologischen Paraphasien, 11 Patienten). Es wurde auditives lexikalisches Entscheiden und Reimentscheiden getestet, bei denen der Prime und der Zielstimulus unrelatiert oder phonologisch relatiert waren (d.h. entweder identisch oder gereimt). Die Probanden mussten zunächst entscheiden, ob der zweite Stimulus ein Wort oder ein Nichtwort war. Bei der zweiten Untersuchung mussten sie entscheiden, ob sich zwei Stimuli hintereinander reimten oder nicht. Die Stimuli waren einsilbige englische Wörter, die zu fünf verschiedenen Prime-Typen gepaart wurden:

- identisch (*book-book*, *chot-chot*)
- Reime (*look-book*, *hot-chot*)
- Nichtwortreime (*zook-book*, *mot-chot*)
- nicht reimende Wörter (*raise-book*, *room-chot*)
- nicht reimende Nichtwörter (*leathe-book*, *yole-chot*)

Zusätzlich gab es noch Füller, um die gleiche Anzahl an Wort- und Nichtwortzielitems zu gewährleisten.

Die lexikalisch-phonologische Verarbeitung war bei Aphasikern mit flüssiger Aphasie gestörter als bei Aphasikern mit nicht flüssiger Aphasie. Aphasiker mit flüssiger Aphasie hatten Zugriff auf sublexikalische phonologische Repräsentationen (konnten Reime beurteilen), hatten aber Probleme, diese phonologische Informationen zu benutzen und Zugriff aufs Lexikon zu erlangen. Flüssige Aphasiker zeigten keinen Reimpriming-Effekt, während Aphasiker mit nicht flüssiger Aphasie sehr wohl einen Reimpriming-Effekt zeigten (Gordon und Baum 1994).

Diese Untersuchungen zeigten, dass auch bei aphasischen Patienten eine Beziehung der phonologischen Ähnlichkeiten im Lexikon besteht. Die Verarbeitung von Reimen

geschieht aber trotzdem anders als bei unbeeinträchtigten Personen (siehe Kapitel 3.3 *Reime bei sprachgesunden Erwachsenen*), ebenso wie es Unterschiede innerhalb der unterschiedlichen Aphasietypen gibt.

Blumstein et al. (2000) meinen, dass der Unterschied in den Ergebnissen der beiden Studien (Blumstein et al. 2000; Gordon und Baum 1994) vielleicht darin liegt, dass die Unterteilungen in Broca- vs. Wernicke-Aphasie bzw. flüssige vs. nicht flüssige Aphasie getroffen wurden.

Nun möchte ich gerne auf die Produktion von Reimen bzw. ihren Einsatz in Therapiestudien eingehen. Meistens werden reimende Wörter zur Therapie von Wortfindungsstörungen eingesetzt:

„Word-finding deficits are one of the most common characteristics of aphasia [...], and may be the only obvious symptom in cases of anomic aphasia“ (Wisenburn und Mahoney 2009, S. 1338).

Wisensburn und Mahoney (2009) führten eine Metaanalyse über Therapien zu Wortfindungsstörungen durch. Sie zitieren Nickels und Best (1996, in: Wisensburn und Mahoney 2009), die schreiben, dass sowohl semantische, als auch phonologische Therapien hilfreich sind bei Wortfindungsstörungen, wobei semantische Therapien aber zu mehr Generalisierung auf ungeübte Items führen. Im Rahmen ihrer Metaanalyse fanden sie nur eine Studie, die auch mit Reimen arbeitete: Raymer et al. (1993) führten eine phonologische Therapie bei „four subjects with severe word-retrieval difficulties stemming in part from failure to activate representations in the phonological output lexicon by way of lexical-semantics“ (Raymer et al. 1993, S. 31). Die Behandlung sollte den Zugriff von der Semantik auf die phonologischen Repräsentationen verbessern, und Generalisierungen von trainierten auf untrainierte Items, sowie von mündlicher auf schriftliche Benennung überprüfen.

Die Stimuli bestanden aus zwei Sets von 30 monosyllabischen abbildbaren Nomen, sowie dementsprechende Bilder und Wörterkärtchen mit diesen Nomen. Jedes Set beinhaltete 10 Trainingswörter, 10 Reimwörter und 10 semantische ähnliche Wörter. In der Baseline und der Überprüfung auf Generalisierung wurden diese Stimuli in drei verschiedenen Aufgaben getestet: mündliches Bildbenennen, schriftliches Bildbenennen und mündliches Lesen des Wortkärtchens.

Die Therapie wurde so aufgebaut, dass der Patient das Bild einer der Trainingswörter benennen sollte. Die Reimwörter und semantisch ähnliche Wörter wurden

hier nicht geübt. Falls die Benennung nicht funktionierte, wurde ein dazu passendes Reimwort präsentiert. Wenn dies auch nicht half, wurde ihm das erste Phonem genannt. Falls es dem Patienten dann noch immer nicht möglich war, das Wort zu produzieren, wurde es ihm vorgesprochen. Danach (egal, bis zu welcher Stufe ihm geholfen wurde) wiederholte der Patient das Wort fünfmal, schließlich noch einmal spontan.

Da die Patienten eine relativ inhomogene Gruppe bezüglich ihrer Fähigkeiten darstellten, zeigte jeder Patient oft ein anderes Muster bezüglich Erwerb und Generalisierung durch die Therapie.

- Bei den Trainingswörtern war die Baseline bei allen vier Patienten ziemlich stabil (verbesserte sich nicht ohne Therapie). Nach Einsetzen der Therapie zeigte sich bei allen vier Patienten ein Erwerb der trainierten Items. Es zeigte sich kein großer Unterschied zwischen nicht trainierten semantisch oder phonologisch ähnlichen Stimuli.
- Bei zwei der vier Patienten zeigte sich eine Generalisierung im mündlichen Benennen auf die phonologisch und semantisch ähnlichen Stimuli.
- Eine Generalisierung auf die schriftliche Wortproduktion zeigte sich auch nur bei zwei von vier Patienten.
- Generalisierungen beim mündlichen Lesen zeigte sich nach Therapie der dement-sprechenden Bilder bei drei von vier Patienten.

Raymer et al. (1993, S. 48) schreiben, dass phonologische Therapien das mündliche Benennen bei Aphasikern verbessern können, obwohl „it may be that other treatments are equally effective, or perhaps even more effective“.

Eine weitere Studie fasst phonologische mit semantischen Faktoren in einer Therapie zusammen: Spencer et al. (2000) stellen in ihrer Studie das Modell einer Reimtherapie in semantischen Kategorien (*Semantic Category Rhyme Therapy*) vor, die sie an einer 47-jährigen aphasischen Patientin (NR) durchführen. Bei der Aphasikerin wurden Beeinträchtigungen in den phonologischen und graphematischen Output-Lexika bei einer relativ intakten semantischen Komponente diagnostiziert (siehe Abbildung 14 *Das Logogenmodell*, auf Seite 48; siehe auch Kapitel 4.2 *Der Zugriff aufs Lexikon*). Auch die phonologische und graphematische Rezeption schien größtenteils intakt

zu sein, sodass die Defizite tatsächlich hauptsächlich am Output lagen. Sie hatte Schwierigkeiten damit, die richtigen Wörter zu produzieren. Spencer und Kollegen entwickelten die Reimtherapie, um die Sprachproduktion der Patientin zu verbessern. Ihre Annahme war, dass „activation would spread from the rhymed word to other entries in the speech output lexicon having similar sounds [...], thus facilitating lexical retrieval.“ (Spencer et al. 2000, S. 568). Zusätzliche semantische Informationen würden den Abruf noch weiter begünstigen.

Vor der Therapie wurden umfangreiche Testungen mit NR durchgeführt, darin begriffen waren auch Reimaufgaben. Sie musste beispielsweise entscheiden, ob sich zwei Wörter reimten. NR verließ sich bei der Lösung der Aufgaben stark auf die graphematische Information der Wörter, wenn sie Schwierigkeiten mit der Verarbeitung der phonologischen Eigenschaften hatte. Diese Aufgaben löste sie zu 75-88% richtig. Bei einer anderen Variante der Aufgabe wurden ihr nur Bilder gezeigt: Wenn sie bei zwei Bildern bestimmen musste, ob sich die Abbildungen reimten, hatte sie 35% richtig, bei vier Bildern nur noch 0%.

Die Reimtherapie in semantischen Kategorien wurde nun folgendermaßen durchgeführt: Spencer und Kollegen erstellten vier Listen mit reimenden Wortpaaren, wobei das zweite Wort das zu produzierende Zielwort darstellte. Das Zielwort gehörte einer von vier semantischen Kategorien an (vierbeinige Tiere, Haushaltsgegenstände, Tischlerwerkzeug und Kleidungsstücke), und das erste Wort des Wortpaares diente lediglich dazu, mit dem Zielwort zu reimen. Jede Liste enthielt Übungsitems, sowie nicht geübte Items, um Generalisierungen zu überprüfen. Die letzte Kategorie (Kleidungsstücke) wurde gar nicht geübt, sondern sollte einen Generalisierungseffekt auf einen anderen semantischen Bereich überprüfen. Am Anfang der Trainingseinheit wurde der Patientin gesagt, mit welcher Kategorie nun gearbeitet werden würde. Erst wenn NR die eine Kategorie gut beherrschte, wurde zur nächsten übergegangen.

Ein konkretes Beispiel wäre hierzu das Wortpaar *course* - *horse* (Spencer et al. 2000, S. 572), die Instruktion beim mündlichen Benennen lautete: „Name a four-legged animal that rhymes with course“ (Spencer et al. 2000, S. 572). Nach richtiger Antwort von NR wiederholte sie das Reimpaar. Dann wurde ihr das Zielwort geschrieben präsentiert und sie sollte es kopieren. Schließlich sollte sie das Reimpaar noch einmal laut aussprechen. Sollte NR zunächst nicht die richtige Antwort geben können, wird ihr der erste Laut des Zielwortes genannt (zum Beispiel: „Das Wort beginnt mit /h/“), wenn es noch immer nicht möglich war, dann wurde ihr der erste

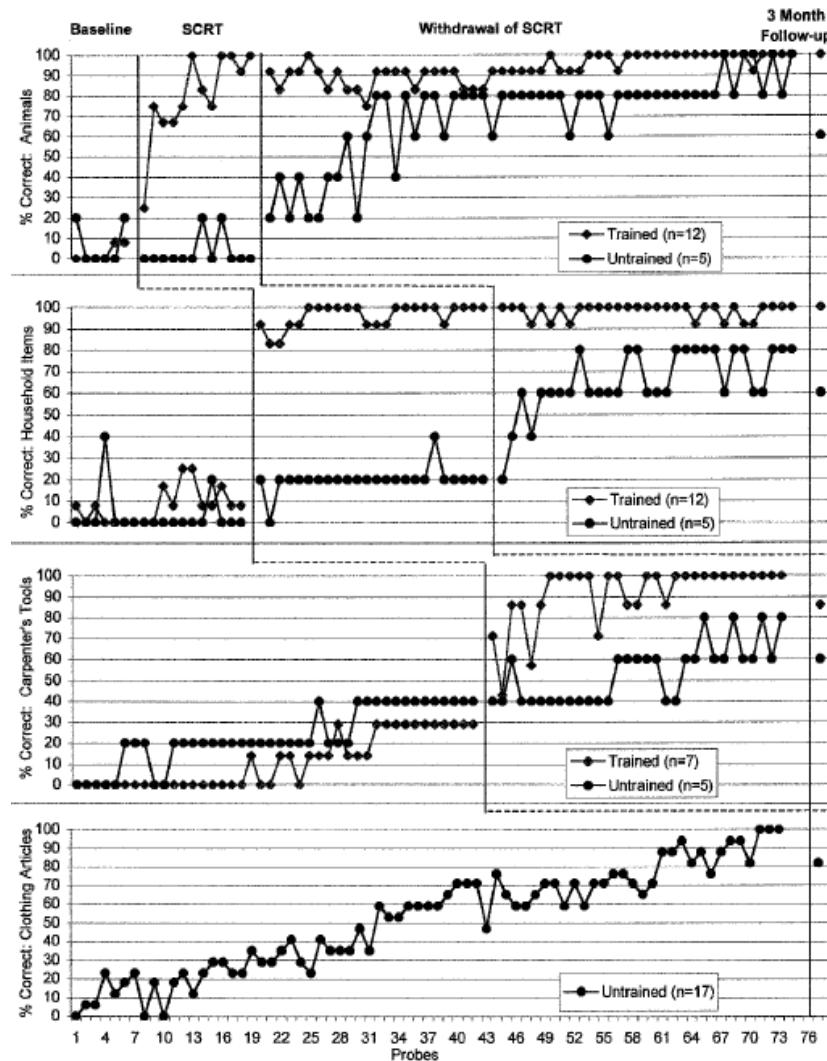


Abb. 12.: Anzahl der korrekten mündlichen Benennungen der trainierten und untrainierten Items in Prozent (Spencer et al. 2000, S. 575)

Laut schriftlich präsentiert. Falls sie es noch nicht benennen konnte, wurde ihr das ganze Wort in beiden Modalitäten präsentiert. Die Ergebnisse nach sieben Monaten Therapie zeigen, dass sich die Kommunikationsfähigkeit der Patientin deutlich verbessert hat, und dies schriftlich sowie mündlich. Die geübten Wörter konnten richtiger produziert werden und es gab einen Generalisierungseffekt auf untrainierte Items in allen vier semantischen Kategorien. Die vierte Kategorie (Kleidungsstücke), die gar nicht geübt wurde, verbesserte sich ebenso, siehe dazu Abbildung 12, die den Fortschritt im mündlichen Benennen darstellt. Die Fortschritte im schriftlichen Benennen verhalten sich ähnlich, obwohl schriftlich weniger geübt wurde als mündlich.

Wie man anhand der Grafik gut erkennen kann, hat die Verbesserung erst eingesetzt, als mit der Reimtherapie in semantischen Kategorien begonnen wurde. Die Verbesserung beim Schreiben erklären sich Spencer und Kollegen dadurch, dass mit Hilfe der Verbesserungen im phonologischen Output-Lexikon auf das graphematische Output-Lexikon leichter zugegriffen werden konnte (siehe dazu Unterkapitel 4.3.6 *Die schriftliche Sprachproduktion*). Spencer und Kollegen befürworteten ein interaktives Sprachproduktionsmodell, da sie bei der Suche nach einem passenden Reim in einer bestimmten semantischen Kategorie ein Feedback von der phonologischen Ebene auf die Semantik annehmen. Somit konnte man sehen, dass die Arbeit mit Reimen die Wortproduktion deutlich verbessern kann. Nach den sieben Monaten wurden die Tests wiederholt, die auch vor der Therapie durchgeführt wurden, und es hat sich gezeigt, dass NRs Leistungen auf andere Bereiche außerhalb der vier semantischen Kategorien übertragbar war (Nichtwörter wiederholen, Maß für richtige, relevante und informative Wörter in zusammenhängender Sprache, Test of Adult Word Finding usw.; Spencer et al. 2000).

3.5. Neurologische Verarbeitung von Reimen

In diesem Kapitel möchte ich auf die neurologische Verarbeitung von Reimen eingehen. Das meistbenutzte Verfahren zur Untersuchung von Reimverarbeitung ist das Elektroenzephalogramm (EEG). Dabei wird mittels Elektroden die elektrische Aktivierung verschiedener Bereiche des Gehirns an der Kopfhaut gemessen (Eysenck und Keane ⁶2010). Diese Aktivierung wird Potential genannt. Da bei der Messung unbeabsichtigt auch andere Dinge (Artefakte wie spontane Hirnaktivität oder Hintergrundverarbeitung des Gehirns) gemessen werden und dadurch die Verarbeitung

des Stimulus' häufig überdeckt wird, wird der gleiche Stimulus mehrmals hintereinander präsentiert. Die einzelnen Verarbeitungen des Stimulus werden extrahiert und (gemessen anhand des Stimulus-Onsets) „übereinandergelegt“, wodurch ein Durchschnitt errechnet wird. Das Ergebnis wird als *ereigniskorreliertes Potential* bezeichnet (EKP, engl.: *event-related potential*, ERP; Eysenck und Keane ⁶2010). Diese Methode kann die Hirnaktivität auf die Millisekunde genau berechnen und wird als Wellenform dargestellt. Unterschieden werden dabei positive und negative Gipfel innerhalb dieser Wellenform.

Die Verarbeitung von Reimen wurde in mehreren Studien untersucht (Coch et al. 2008, Grossi et al. 2001, Rugg 1984b, Rugg und Barrett 1987, Wagensveld et al. 2012). Dabei wurden Wortpaare (visuell oder auditiv) präsentiert und die Probanden sollten entscheiden, ob sich die Wortpaare reimten oder nicht. Weber-Fox et al. (2003) beschreiben den Vorgang, der hierbei vonnöten ist, folgendermaßen:

The abilities necessary to make rhyming judgments for pairs of words (i.e., prime-target) include retrieving the phonological representation of a prime, holding it in working memory, and segmenting it into its onset and rime constituents. The same processes would then be utilized for processing the target word in the pair. To complete the rhyme judgment, the rime constituent of the prime must be compared to the rime of the target.
(Weber-Fox et al. 2003, S. 128)

Die Reimverarbeitung zeigt sich als negative Krümmung bei Nichtreimung im Vergleich zu Reimung ca. 450ms nach Stimulusonset, hauptsächlich bei mittigen und rechtshemisphärischen Elektroden (Coch et al. 2008, Grossi et al. 2001, Rugg 1984b, Rugg und Barrett 1987). Dies gilt auch für Nichtwörter (Rugg 1984a, in: Wagensveld et al. 2012), daher kann dieser Effekt als phonologisch, und nicht als semantisch angesehen werden (Coch et al. 2008). Coch et al. (2008) schreiben auch, dass diese Ergebnisse ungewöhnlich sind, da bisher angenommen wurde, dass die rechte Hemisphäre eher weniger mit phonologischer Verarbeitung zu tun hätte. Kürzlich hätte sich allerdings gezeigt, dass bei Reimaufgaben die rechte Hemisphäre aktiviert wird, wenn auch etwas später als die linke.

Die Reimverarbeitungs-komponente wird als N450, manchmal auch als N400 oder N350 (vgl. Weber-Fox et al. 2003) bezeichnet. In nachfolgenden Studien hat man herausgefunden, dass der N450-Reimeffekt bei Kindern und Erwachsenen, bei auditiver

und visueller Präsentation der Stimuli, sowie bei simplen (isolierten Graphemen) und komplexen Stimuli (mehrsilbigen Nichtwörtern) zu finden ist (Wagensveld et al. 2012).

Wagensveld et al. (2012) stellten die Frage, ob dieser N450-Reimeffekt tatsächlich nur die Reimverarbeitung widerspiegelt, oder gar einen allgemeinen Effekt für globale Ähnlichkeit darstellt (*global similarity effect*).

Der globale Ähnlichkeitseffekt beschreibt die Tatsache, dass es schwieriger ist, zu entscheiden, ob sich zwei Wörter reimen oder nicht, wenn sie sich in großem Maße ähneln oder überlappen (*phonological overlap*). Ein Beispiel hierfür wären die Wörter *ball* und *bell*. Die Entscheidung der Probanden dauert dabei länger und führt zu mehr Fehlern als beispielsweise bei der Entscheidung zwischen *ball* und *fall*. Wagensveld et al. (2012) berichten von Experimenten mit Kindern und Erwachsenen, die diesen globalen Ähnlichkeitseffekt untersuchen sollten. Dabei wurden Kindern und Erwachsenen immer drei Silben (*bis-diz-bun*) präsentiert, wobei immer zwei Silben das gleiche Anfangsphonem (*bis-bun*) und zwei Silben drei ähnliche Phoneme besaßen (*bis-diz*). Die Probanden sollten die zwei Wörter zusammentun, die ihrer Meinung nach zusammengehörten. Kinder wählten dabei öfter das Silbenpaar, das die meisten ähnlichen Phoneme besaß, während Erwachsene öfter das Silbenpaar wählten, das das gleiche Anfangsphonem besaß. Somit kann man sagen, dass die Wichtigkeit phonematischer Eigenschaften und Informationen im Alter zunimmt (Wagensveld et al. 2012; siehe auch Kapitel 3.2 *Die phonologische Bewusstheit*).

Wagensveld et al. (2012) führen verschiedene Studien heran, die berichten, dass das Ausmaß der N450-Komponente steigt, je mehr phonologische Überlappung die Stimuli aufweisen. Dies geht von der Alliteration bzw. von dem gleichen letzten Phonem, über die Reimkomponente bis hin zur Silbe und gilt für Wörter, wie auch für Nichtwörter.

Um das Ausmaß der globalen Ähnlichkeit im Hinblick auf den N450-Reimeffekt zu untersuchen, führten Wagensveld et al. (2012) ein Experiment an 20 Erwachsenen durch. Den Probanden wurden Wortpaare mit phonologischer Überlappung (*huis-haas*, dt. Haus-Hase), Reimpaare (*kaas-haas*, dt. Käse-Hase) und unrelatierte Wortpaare (*dun-haas*, dt. dünn-Hase) auditiv präsentiert. Sie mussten entscheiden, ob sich die beiden Wörter reimten oder nicht, es gab somit Ergebnisse in Form von Verhaltensmessungen (Ja-/Nein - Entscheidung) und in Form von EEG-Messungen.

Die Ergebnisse zeigten einen globalen Ähnlichkeitseffekt bei der Verhaltensmessung:

Wortpaare mit phonologischer Überlappung wurden von den erwachsenen Probanden langsamer beurteilt als Reimpaare und unrelatierte Wortpaare. Dadurch stellten Wagensveld und Kollegen fest, dass der globale Ähnlichkeitseffekt nicht nur ein Entwicklungsprozess ist, der im Kindesalter auftritt. Die EEG-Messungen zeigten einen N450-Reimeffekt bei Reimpaaren im Vergleich zu nicht-reimenden Wortpaaren zwischen 250 und 500 Millisekunden, eingeschlossen der Wortpaare mit phonologischer Überlappung. Somit kann man sagen: „The neural data [...] did not show a global similarity effect“ (Wagensveld et al. 2012, S. 65): die N450-Komponente präsentiert unterschiedliche Daten bei Reimpaaren und Wortpaaren mit phonologischer Überlappung, während unrelatierte und phonologisch ähnliche Wortpaare beim frühen lexikalischen Zugriff vermutlich gleichermaßen verarbeitet werden (vgl. Abbildung 13). Bereits Rugg (1984b) vermutete, dass geprimte Reimwörter Wortkandidaten aktivieren, die eine mögliche Reimung darstellen würden. Wenn dann als zweiter Stimulus ein Wort auftritt, das sich mit dem Prime nicht reimt, würde der N450-Reimeffekt ausgelöst werden.

Coch et al. (2005) führten ein Experiment zum Reimentscheiden bei Kindern (im Alter von 6 bis 8 Jahren) durch. Der gleiche N450-Reimeffekt hat sich gezeigt, allerdings verzögerte sich der Onset des Reimeffekts bei denjenigen, die bei Messungen zur phonologischen Bewusstheit generell schlechter abschnitten.

In engem Zusammenhang mit dem N450-Reimeffekt steht die N400-Komponente, die bei semantischen Inkongruenzen auftritt (um Verwechslungen vorzubeugen wird der Reimeffekt daher häufig als N450-Effekt bezeichnet). Dies gilt für sinnngemäße Unstimmigkeiten, wie auch für Änderungen in der Wortform (Wagensveld et al. 2012; Rugg 1984b). Rugg und Barrett (1987) meinen, dass die N400- und N450-Komponenten gleiche zugrundeliegende Prozesse innehaben. Das Verständnis von der Reimverarbeitung im Zusammenhang mit anderen phonologischen Eigenschaften wird noch weiterer Erforschung bedürfen, zumal sich auch gezeigt hat, dass die Gewinnung der Ergebnisse stark von der Instruktion der Aufgabe abhängt. Zum Beispiel zeigt sich beim passiven Hören von Reimstimuli kein N450-Effekt (Perrin und Garcia-Larrea 2003, in: Wagensveld et al. 2012), andere Autoren baten die Probanden um Lexikalitätsurteile bei der Präsentation von Reimstimuli (weiterführende Literatur in Wagensveld et al. 2012). Bei Wagensveld und Kollegen hingegen wurde darum gebeten, die Reimung zweier Wörter zu entscheiden, und dies führte zum N450-Effekt in beiden Bedingungen, bei der phonologischen Überlappung oder bei

3. Der Forschungsstand zu Reimen in der Psycho-, Patho- und Neurolinguistik

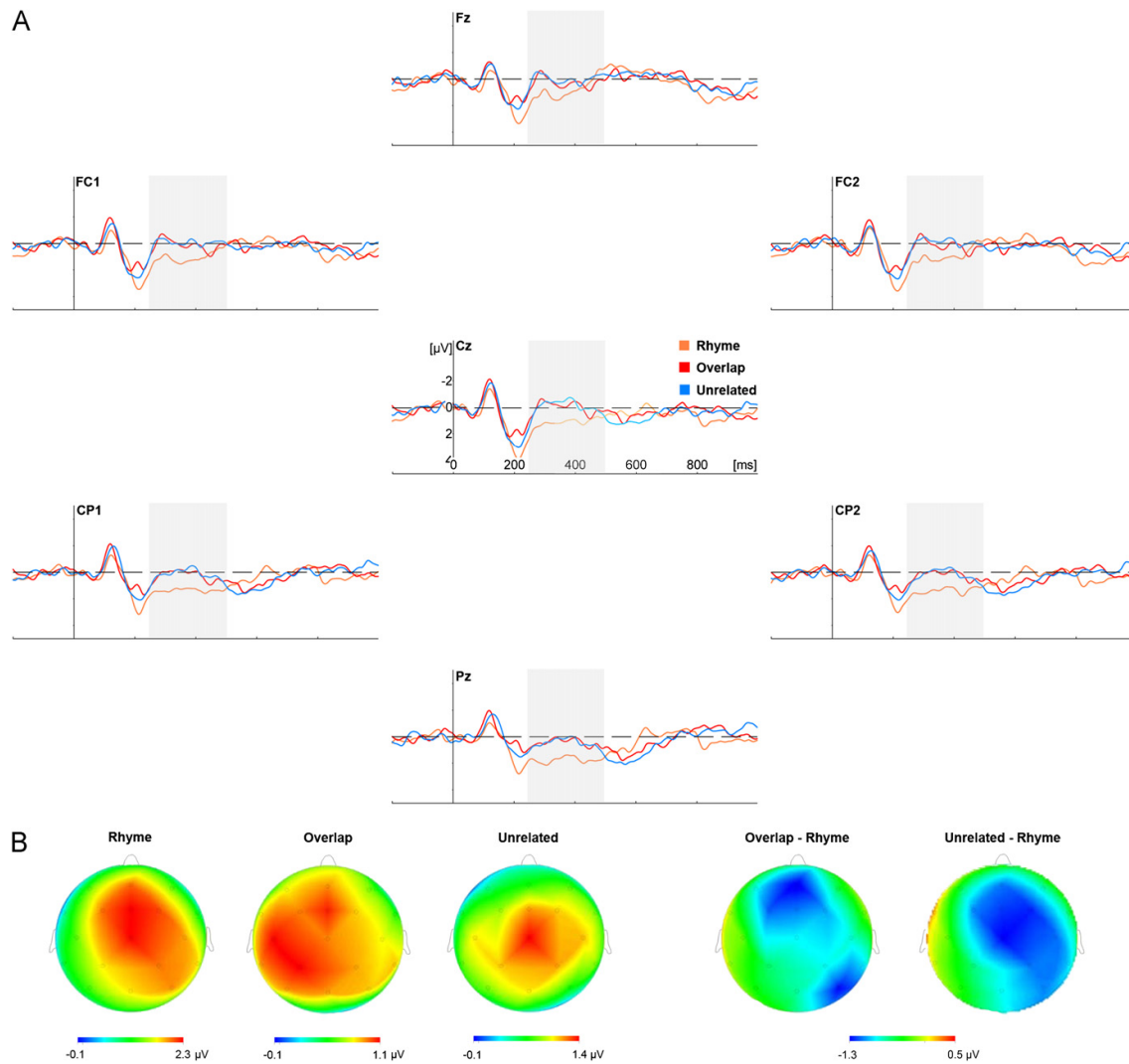


Abb. 13.: (A) Übersicht der EKPs bei allen drei Bedingungen. Die Negativität ist nach oben gerichtet. (B) Topographische Karten der drei Bedingungen und Differenzdarstellung (Wagensveld et al. 2012, S. 66)

den unrelatierten Wortpaaren.

Ein weiterer wichtiger Untersuchungsgegenstand ist bei der Reimverarbeitung das Verhältnis zwischen Graphematik und Phonologie. Rugg und Barrett (1987) meinten bezüglich des Reimentscheidens bei schriftlichen Aufgaben, dass „die Orthographie [...] einen kraftvollen Einfluss auf die Effizienz [hat], mit der solche Entscheidungen getroffen werden“ (S. 337). Der Einfluss der Graphematik auf die Reimentscheidung ist auch gegeben, wenn die Wörter nur gehört werden (Rugg und Barrett 1987). In ihrer Studie untersuchen Rugg und Barrett (1987) das Verhalten und die EKPs bei der Reimentscheidung bei visuell präsentierten sich reimenden und nicht reimenden orthographisch ähnlichen und nicht ähnlichen Wortpaaren. Es hat sich gezeigt, dass die Geschwindigkeit und Richtigkeit der Antworten stark von der orthographischen Ähnlichkeit beeinflusst werden. Die Entscheidungen fielen leichter, wenn die reimenden Wörter graphematisch ähnlich waren, und schwieriger, wenn sie graphematisch ähnlich waren, sich aber nicht reimten. Zu dem gleichen Schluss kamen auch Weber-Fox et al. (2003). Rugg und Barrett erklären diese Erscheinung dadurch, dass die Eigenschaften des ersten Wortes eine gewisse Erwartungshaltung gegenüber dem zweiten Wort entwickeln, möglicherweise handelt es sich gar um graphematisches Priming, bei dem Reim- bzw. Wortkandidaten kodiert werden, die sich phonologisch reimen würden (bei dem Wort *pale* beispielsweise wären dies die Reimkomponenten *-ale* bzw. *-ail*). Die EKPs zeigten die bereits bekannte negative Komponente bei ca. 400-450 Millisekunden (N450-Reimeffekt).

Coch et al. (2008) führten ebenfalls eine visuelle Reimaufgabe durch. Es wurden Wortpaare, Nichtwortpaare und einzelne Buchstaben getrennt präsentiert und die erwachsenen Probanden sollten entscheiden, ob sich die jeweiligen Paare reimen. Es handelt sich dabei um eine der ersten Studien, die Buchstabenreime testen (z.B.: *b* reimt sich mit *c*, aber nicht mit *a*). Die Autoren halten eine solche Aufgabe für wichtig, da „letter rhyme tasks simultaneously index the two best predictors of ease of learning to read: letter name knowledge and phonological awareness“ (Coch et al. 2008, S. 230). Auch hier hat sich der N450-Reimeffekt gezeigt (zwischen 230 und 500 Millisekunden), und zwar bei allen drei Bedingungen. Bei den Buchstaben war der Effekt allerdings am geringsten. Die Autoren meinen, es handelt sich womöglich um einen phonologischen Primingeffekt. Die visuelle N450-Komponente scheint ungeachtet der Menge orthographischer Information (von einzelnen Buchstaben bis hin zu Nichtwörtern) hervorgerufen zu werden.

Grossi et al. 2001 führten ein ähnliches, visuelles Experiment mit Probanden in acht Altersgruppen durch (im Alter von 7 bis 23 Jahren). In allen Altersgruppen zeigte sich der N450-Reimeffekt, bei den jüngeren Probanden gar bei reimenden Wortpaaren, wenn auch in geringerem Ausmaß als bei nichtreimenden Wortpaaren. Bei Probanden ab 18 Jahren lösten nur die nicht reimenden Stimuli den N450-Reimeffekt aus. Das Einsetzen des Reimeffekts ändert sich somit im Laufe der Entwicklung nicht. Auch Weber-Fox et al. (2003) postulieren, dass die Prozesse, die bei der Reimentscheidung notwendig sind, bereits sehr früh entwickelt sind.

Es sind weitere Forschungen notwendig, um die Bandbreite der Verarbeitung von Reimen genauer zu untersuchen, besonders den Zusammenhang zwischen Schriftsprache und Phonologie, sowie möglicherweise durch Untersuchungen mit Kindern vor dem Schriftspracherwerb.

4. Theoretischer Hintergrund

Dieses Kapitel soll das Lexikon, dessen Aufbau und den Abruf von Wörtern aus dem Lexikon zum Thema haben. Dies ist wichtig, um die Analyse der Ergebnisse des in dieser Arbeit durchgeführten Experiments zu analysieren und zu deuten. Es soll um die Faktoren gehen, die die Wortproduktion und Wortwahl beeinflussen. Nach einer kurzen Einleitung werden bekannte Sprach- bzw. Wortproduktionsmodelle vorgestellt und verglichen. Zunächst versuche ich, das Lexikon an sich zu umreißen, um dann auf den Zugriff auf Wörter aus diesem Lexikon einzugehen.

4.1. Das mentale Lexikon

Das mentale Lexikon kann folgendermaßen definiert werden:

„Psychologische Repräsentation lexikalischer Informationen, Sammelname für die mentale Organisation des Wortvorrates.“ (Glück ⁴2010, Stichwort *Mentales Lexikon*, S. 420)

„Bezeichnung für den mental organisierten und repräsentierten Wortschatz, auf den in der Sprachverarbeitung zugegriffen wird (‘lexikalischer Zugriff’)“ (Bußmann ⁴2008, Stichwort *Mentales Lexikon*, S. 430)

Das mentale Lexikon ist laut Schwarz (1996, in: Vater ⁴2002, S. 207) „der Teil des Langzeitgedächtnisses, in dem die Wörter einer Sprache mental repräsentiert sind“.

Libben und Jarema (2002, S. 2) bezeichnen das mentale Lexikon als „das Rückgrat der Sprachfähigkeit“:

„It is the theoretical characterization of lexical knowledge and encompasses research on how that knowledge is organized to function effectively in the millisecond time domain without conscious intervention, to self-organize as a result of new experience, and to maintain information integrity over a lifetime.“ (Libben und Jarema 2002, S. 2)

Die Untersuchungen im Zusammenhang mit dem mentalen Lexikon zeigen, dass es Generalisierungen über alle Sprachen hinweg geben könnte, wodurch das mentale Lexikon auch als „menschliche mentale Organisation“ (Libben und Jarema 2002, S. 3) bezeichnet werden könnte. Jedoch ist es auch möglich, dass aufgrund der unterschiedlichen Sprachen und ihrer Charakteristika das mentale Lexikon unterschiedlich strukturiert ist (Libben und Jarema 2002). Ein unbeeinträchtigtes mentales Lexikon enthält den (gebräuchlichen) Gesamtwortschatz einer Sprache, wobei der aktive Wortschatz einen Teil des passiven Wortschatzes darstellt (Vater ⁴2002, S. 207). Doch wie und nach welchem Muster die Wörter im mentalen Lexikon angeordnet sind, ist nicht ganz eindeutig zu bestimmen. Neben Experimenten mit unbeeinträchtigten L1-Sprechern einer Sprache werden noch andere Probandengruppen genauer untersucht, um über den Aufbau des Lexikons und seinen Subkomponenten mehr zu erfahren: Zum Einen werden Sprachstörungen als ein Indiz dafür herangezogen, wie das mentale Lexikon strukturiert ist (Vater ⁴2002), zum Anderen werden dazu der Spracherwerb und der Bilingualismus untersucht (Libben und Jarema 2002).

Ein wichtiger Begriff im Zusammenhang mit dem mentalen Lexikon ist der des *Lexems*. Ein Lexem ist ein lexikalisches Wort, das eine Einheit im Lexikon darstellt. Cholewa und Mantey (2007, S. 152) verstehen unter einem Lexem „abstrakte lexikalische Repräsentationen von Wörtern [...], die noch nicht hinsichtlich ihrer Flexionsmerkmale ausgeformt sind. [...] Das Grundwort selbst verfügt noch nicht über eine aussprechbare Form und ist in Bezug auf grammatische Kategorien nicht festgelegt“. Das heißt, dass die Wörter *Tisch* (Nominativ Singular) und *Tische* (Nominativ Plural) verschiedene Formen desselben Lexems darstellen (in Kapitel 4.3 *Sprachproduktionsmodelle* sehen wir, dass auf der Ebene der Repräsentationen der Lexeme die syntaktischen Eigenschaften kodiert werden).

Homonyme sind beispielsweise verschiedene Lexeme: *Bank* (Institut) und *Bank* (Sitzbank), während flektierte Wörter häufig nicht als eigenständige Lexeme angesehen werden, sondern als Ableitungen ihrer Grundstämme (*Bänke*) (Vater ⁴2002). „Flektierte Wörter wie *leb-t-e* [werden] als Kombination separat gespeicherter Morpheme (Wurzeln/Affixe) aufgefasst [...] ; irreguläre Formen wie *war* oder *wurde* werden nicht dekomponiert, sondern als Ganzes gespeichert“ (Vater ⁴2002, S. 208).

Die meisten Autoren sind sich darüber einig, dass das Lexikon aus verschiedenen unabhängigen Sublexika für phonologische, orthographische, und syntaktische Repräsentationen besteht, die aber doch miteinander verbunden sind. Zusätzlich dazu

gibt es noch eine semantische Subkomponente, in der die Wörter nach semantischen Merkmalen angeordnet sind (Vater ⁴2002), sowie je nach Theorie weitere Sublexika nach morphologischen oder beispielsweise prosodischen Merkmalen (vgl. Kapitel 4.3 *Sprachproduktionsmodelle*).

Nun wird zuerst ein Modell vorgestellt, das nicht nur die Sprachproduktion, sondern auch die Sprachrezeption darstellt. Damit soll eine Variante präsentiert werden, wie man sich das Lexikon bzw. die Sprachverarbeitung vorstellen kann. Es handelt sich hierbei um das *Logogenmodell*. Huber et al. (2006, S. 76) bezeichnen ein Logogen als „die mentale Speichereinheit einer Wortform im Unterschied zur Wortbedeutung [...] Input-Logogene bestimmen den Wortzugriff beim Hören und Lesen, Output-Logogene bestimmen den Wortabruf beim Sprechen und Schreiben“ (man beachte hierbei, dass im weiteren Verlauf der Arbeit auch der Wortabruf beim Sprechen und Schreiben als *Zugriff* bezeichnet wird). Das Logogenmodell beschreibt somit nur die Sprachverarbeitung von einzelnen Wörtern, nicht die von größeren Ebenen wie Sätzen oder Ähnlichem. Das Modell basiert auf dem Einzelfallansatz (im Gegensatz zum Syndromansatz), bei dem die Fähigkeiten von einzelnen sprachbeeinträchtigten Personen untersucht werden. Aufgrund der unterschiedlichen Leistungen von Patienten können die unterschiedlichsten Störungen angenommen werden, wodurch sich schließlich der Aufbau des Logogenmodells entwickelt hat (Stadie et al. 1994). Es gibt viele verschiedene Versionen des Logogenmodells, in Abbildung 14 ist jene Version abgebildet, die auch in LeMo (De Bleser et al. 2004) verwendet wird. Es handelt sich hierbei um ein nicht interaktives Modell.

Wie man anhand der Grafik sehen kann, gibt es hier nicht nur ein mentales Lexikon, sondern vier: das phonologische Input- und Outputlexikon für die mündliche, und das graphematische Input- und Outputlexikon für die schriftliche Wortverarbeitung. Im Zentrum liegt die Semantik, das die Wortbedeutungen repräsentiert. Vor den Lexika bei der Sprachrezeption und nach den Lexika bei der Sprachproduktion befinden sich die jeweiligen sogenannten Buffer, die eine kurzfristige Speicherung des Sprachinputs bzw. -outputs ermöglichen. Man kann anhand der Pfeile nachverfolgen, wie die Sprachverarbeitung vor sich geht, inwieweit die Semantik eingebunden werden muss, oder nicht. Weitere Routen sind die APK (auditiv-phonologische Korrespondenzroute), die GPK (Graphem-Phonem-Konversion), die PGK (Phonem-Graphem-Konversion) und die PRS (phonologische Rückkoppelungsschleife, sie ermöglicht durch „internes Sprechen eine wiederholte Reaktivierung der Repräsenta-

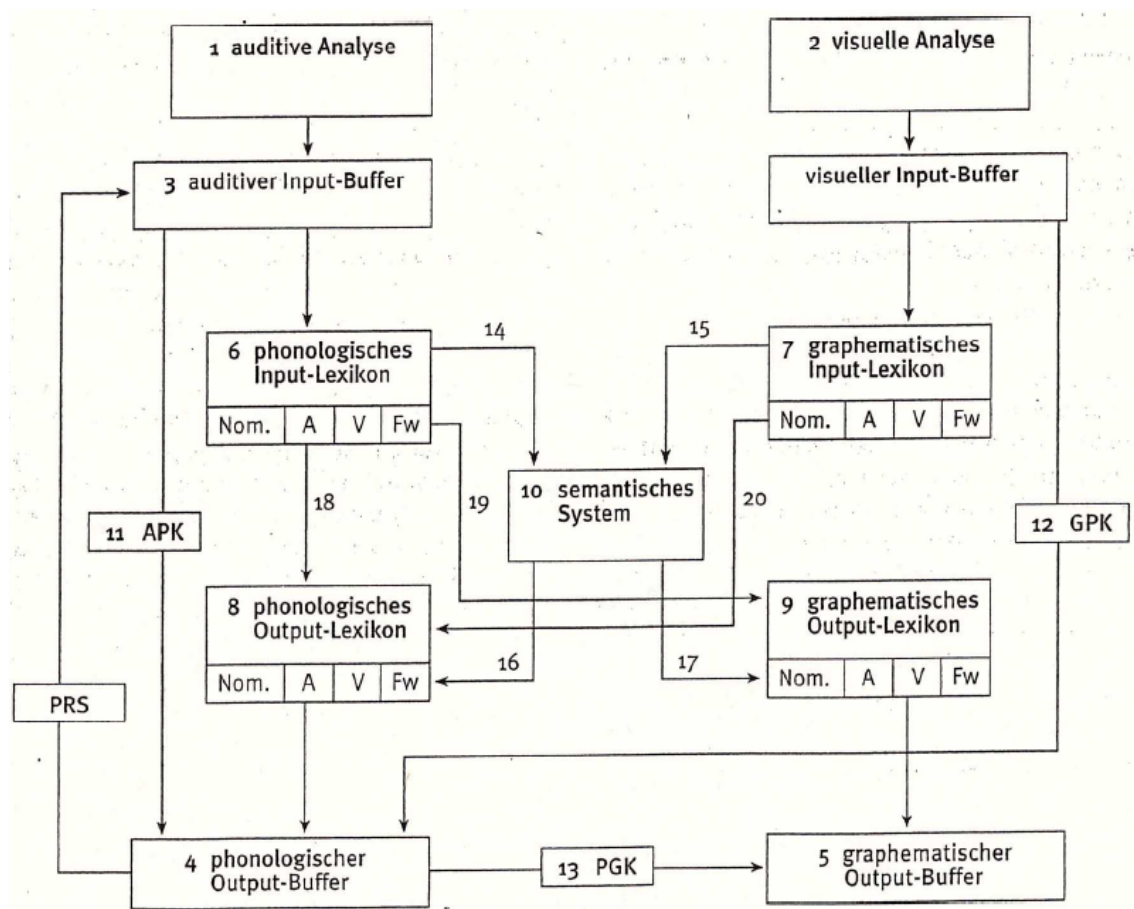


Abb. 14.: Das Logogenmodell in LeMo in Anlehnung an Patterson 1988 (De Bleser et al. 2004, S. 7)

tionen in den Buffersystemen“ (De Bleser et al. 2004, S. 9).

Das Logogenmodell wurde jetzt bereits vorgestellt, da es einerseits, wie gesagt, das gesamte Spektrum der Sprachverarbeitung umfasst, andererseits aber nicht - oder zumindest nicht so wie die später in dieser Arbeit behandelten Sprachproduktionsmodelle - genauer erklärt, wie Sprachproduktion vor sich gehen könnte. Zusätzlich umfasst es nur die Verarbeitung von einzelnen Wörtern, jedoch ist bereits zu erkennen, wie die graphematische Sprachverarbeitung im Vergleich zur phonologischen veranschaulicht wird (vgl. Kapitel 4.3.6 *Die schriftliche Sprachproduktion*).

Nun wird auf die Sprachproduktion (der Zugriff auf das Lexikon), näher eingegangen (im Logogenmodell umfasst dies den unteren Teil des Modells ab den Output-

Lexika mit oder ohne Einwirkung der Semantik. Statt „mentales Lexikon“ wird im Folgenden meist einfach nur „Lexikon“ verwendet.

4.2. Der Zugriff aufs Lexikon

Das Benennen von Dingen, die Wortwahl, der Ausdruck unserer Gedanken - dies sind sehr komplexe Vorgänge im menschlichen Sprachgebrauch. Für viele Objekte und Ereignisse um uns herum gibt es viele verschiedene Ausdrücke und Beschreibungsvarianten. Um beispielsweise ein Pferd zu benennen, könnte man sagen, es sei ein *Pferd*, oder ein *Tier*, ein *Ross*, vielleicht ein *Schimmel*, bei einem weiblichen Exemplar eine *Stute* - und alle Begriffe könnten eine Beschreibung ein und desselben Lebewesens sein. Was also bringt uns dazu, eine dieser Varianten auszuwählen und zu produzieren - und wieso produzierten wir nicht eine andere? Natürlich hat jeder Begriff andere Teilbedeutungen, passt in einen anderen kontextuellen Rahmen, doch gibt es noch andere Faktoren, die und bei unserer Wortwahl beeinflussen?

Anhand des oben behandelten Logogenmodells könnte man die Sprachproduktion folgendermaßen beschreiben: Es gibt eine Intention, eine Bedeutung, die der Sprecher ausdrücken möchte, dadurch wird auf das Outputlexikon zugegriffen (auf das phonologische bei der mündlichen Produktion, auf das graphematische bei der schriftlichen) und das zur Bedeutung passende Wort ausgewählt, im Lexikon wird die korrekte Wortform ausgewählt (De Bleser et al. 2004, Shelton und Weinrich 1997, Stadie et al. 1994). Im Buffer wird dies gespeichert und dann ausgesprochen bzw. niedergeschrieben. Man kann aber auch ohne Semantik sprechen oder schreiben (zum Beispiel beim Nachsprechen oder beim Schreiben von Nichtwörtern, die keinen Eintrag im Lexikon besitzen).

Es soll die Frage behandelt werden, nach welchen Kriterien die Sprecher einer Sprache Wörter auswählen, um sie auszusprechen oder niederzuschreiben. Der Zugriff aufs Lexikon wurde bisher mithilfe von Personen mit Sprachdefiziten, Versprechern und durch zahlreiche Priming-Experimente (vgl. Kapitel 5.1 *Priming-Experimente*) näher untersucht (Rapp und Samuel 2002).

Alario und Kollegen (2003, S. 372f) definieren den Zugriff aufs Lexikon folgendermaßen:

„Lexical access is the process by which a communicative intention leads to

the spoken or written production of words that convey the corresponding message. Central to this process is the activation and selection of entries in the mental lexicon, as well as the retrieval of the different properties of the words required for production.“

Verschiedene Modelle zur Wortproduktion wurden bereits entwickelt und sollen im Folgenden näher vorgestellt werden. Gemeinsam ist ihnen, dass die Produktion eines bestimmten Wortes von semantischen Kriterien, aber auch nach phonologischen und orthographischen Kriterien geschehen kann, bei manchen Autoren auch nach morphologischen Kriterien oder Betonungsmustern. Rapp und Samuel (2002) schreiben, dass ein Dichter im Rahmen eines Gedichts die Wörter anders wählen wird, als in der gesprochenen Spontansprache - er sucht beispielsweise nach Reimwörtern, die aber auch semantisch zum Inhalt passen.

Alario und Kollegen (2003, S. 373) fassen die Wortproduktion am Beispiel einer Aufgabe zur Bildbenennung zusammen:

- Zuerst findet die Konzeptualisierungsphase statt (*conceptualization stage*), bei der das Objekt identifiziert und dessen Konzept aktiviert wird. Dadurch werden verschiedene lexikalische Einträge aktiviert.
- Dies führt zur Lexikalisierungsphase (*lexicalization stage*). Hier wird das Wort ausgewählt, das am besten zum Objekt passt und seine lexikalischen Eigenschaften werden vom mentalen Lexikon abgerufen.
- Schlussendlich kommt man zur Produktionsphase (*output stage*) und das ausgewählte Wort wird entweder mündlich oder schriftlich produziert.

Hier fällt bereits auf, dass sich die Beispiele mit dem Dichter und der Bildbenennung einander etwas widersprechen - es handelt sich ja auch um unterschiedliche Anforderungen. In der Beschreibung von Alario und Kollegen (2003) scheint es eindeutig zu sein, dass erst NACH Auswahl des Wortes in der Lexikalisierungsphase dessen phonologische Eigenschaften kodiert werden. Rapp und Samuel (2002) bestätigen, dass in der Literatur zu Sprachproduktion die Meinung vorherrscht, semantische Eigenschaften würden vor phonologischen Eigenschaften kodiert werden. Ein Dichter verfährt bei der Suche nach Reimwörtern aber anders: Er stellt bereits bei der Auswahl seiner Wörter Anspruch auf phonologische Eigenschaften oder passt deren Inhalt gar daran an. Nun stellt sich die Frage, die Rapp und Samuel (2002) in

ihrem Artikel auch stellten: Handelt es sich bei der Anwendung phonologischer Faktoren bei der Wortauswahl um Sonderfälle der Sprachverwendung (wie zum Beispiel in der Dichtung) oder kann man von generellen Vorgängen sprechen, die auch in der Spontansprache auftreten?

Diese Fragestellung soll im Folgenden durch Veranschaulichung verschiedener Modelle zur Sprachproduktion beleuchtet werden. Hierbei kann man im Groben serielle Modelle von interaktiven Modellen unterscheiden.

4.3. Sprachproduktionsmodelle

Im Laufe der Zeit haben sich verschiedene Sprachproduktionsmodelle entwickelt, die sich manchmal in grundlegenden Punkten unterscheiden, manchmal nur in Einzelheiten. Ich möchte zunächst versuchen, einen generellen Überblick über Sprachproduktionsmodelle zu geben (was ihnen gemeinsam ist), sowie dann Modelle vorstellen, die sich in wichtigen Punkten unterscheiden, und die die Ergebnisse des Experiment im Rahmen dieser Arbeit aus verschiedenen Blickwinkeln betrachten lassen.

Caramazza (1997, S. 177) schreibt, dass es zwei fundamentale Punkte in der Sprachproduktion gibt, über die sich die Autoren zumeist einig sind:

- Semantische und syntaktische Informationen, sowie Informationen über die lexikalische Form der Spracheinheiten bilden unabhängige Repräsentationsebenen.
- Auf diese Ebenen wird im Rahmen der Sprachproduktion womöglich sequenziell zugegriffen.

Den Vorgang des Zugriffs aufs Lexikon beschreiben Caramazza (1997), Rapp und Goldrick (2000), und Dell (1997) folgendermaßen:

- Im ersten Schritt erfolgt eine Auswahl aus semantisch und syntaktisch spezifizierten lexikalischen Repräsentationen (oder *Lemmata*). Dies kann man als Mapping von konzeptuellen, semantischen Repräsentationen auf die Lemmata verstehen. Hierbei werden die nicht phonologischen Repräsentation eines Wortes kodiert (sondern nur semantische und grammatische Repräsentationen). Dieser Schritt basiert somit auf der Bedeutung.

- Der zweite Schritt stellt den Zugriff auf die phonologischen Repräsentationen dar, und es handelt sich hierbei um das Mapping von den Lemmata auf die einzelnen Wortformen. Es erfolgt eine dementsprechende Auswahl aus phonologischen Repräsentationen (oder *Lexemen*).

Dies sind Faktoren, die den meisten Sprachproduktionsmodellen gemeinsam sind (sie betreffen zumindest die mündliche Sprachproduktion - über die schriftliche Sprachproduktion wird mehr in Kapitel 4.3.6 *Die schriftliche Sprachproduktion* berichtet). Im Übrigen weisen die Modelle viele Unterschiede auf: Fungieren die einzelnen Ebenen eigenständig oder interaktiv, gibt es eine zusätzliche morphologische Repräsentationsebene, sind es ganzheitliche oder aufgeteilte Repräsentationen usw., aber die Unterteilung von Lemma und Lexem zieht sich jedoch durch viele verschiedene Modelle (siehe Caramazza 1997) und wird als *two-stage model of lexical access* bezeichnet. Miceli et al. (1997, S. 36) schreiben aber, dass es nicht ganz klar ist, ob „lexical-semantic information [is] represented at the lemma level or are lemmas abstract nodes whose only function is to link semantic, grammatical, and lexical form information“.

Evidenz für eine solche Unterteilung liefern Vertauschungen: Wortvertauschungen geschehen immer innerhalb der gleichen grammatikalischen Klasse, Lautvertauschungen sind immer phonologisch ähnlich oder geschehen in einer phonologisch ähnlichen Umgebung. Ein anderes Indiz dafür sind Reaktionszeitmessungen, das Muster beim Zögern bei normaler oder aphasischer Sprache, oder das Tip-of-the-tongue-Phänomen (das Wort liegt dem Sprecher „auf der Zunge“, seine Wortform kann aber nicht abgerufen werden), bei dem das Lemma bekannt ist, das Lexem aber nicht. Worüber sich die Autoren häufig nicht einig sind, ist der Aufbau der Lemma-Ebene: Ist sie modalitätsunabhängig, oder bezieht sie sich nur auf den phonologischen oder den orthographischen Output, beinhaltet sie semantische und/oder syntaktische Repräsentationen usw. Sicher ist nur, dass sie keine phonologischen Informationen enthält.

Ein weiteres Modell, das bereits drei unterschiedliche Ebenen postuliert, haben Roelofs (1992), Bock und Levelt (1994), und Jescheniak und Levelt (1994, alle in: Caramazza 1997) entworfen und ist folgendermaßen aufgebaut:

- Es gibt eine konzeptuelle (semantische) Ebene, auf der die Bedeutung eines Wortes durch ein Set von markierten Verbindungen zwischen Konzeptknoten und anderen Knoten im Netzwerk dargestellt wird. Jeder Konzeptknoten ist mit Lemma-Knoten verbunden.

- Dann folgt die Lemma-Ebene. Sie ist modalitätsunabhängig (d.h. es gibt keine Unterscheidung phonologischer oder graphematischer Sprachrepräsentationen) und besteht aus syntaktischen Knoten, die Eigenschaften wie die grammatische Klasse (Nomen, Verb, usw.) oder Genus (maskulin, feminin, neutrum) repräsentieren. Jeder Lemma-Knoten ist mit Lexem-Knoten verbunden.
- Schließlich wird die Lexem-Ebene erreicht. Sie besitzt Verbindungen zu den segmentalen Knoten (Laute, Schriftzeichen) und spezifiziert phonologische und graphematische Wortformen.

„Lexical access in this model is represented by the sequential selection of lemma (and hence the syntactic features that define a word) and lexeme nodes through spreading activation emanating from the lexical concept node.“ (Caramazza 1997, S. 181)

Auch Rapp und Goldrick (2000) schreiben, dass die meisten Sprachproduktionsmodelle drei Hauptrepräsentationen unterscheiden (Semantik, Syntax und Phonologie), wobei die Orthographie häufig außer Acht gelassen wird (mehr dazu unter Kapitel 4.3.6 *Die schriftliche Sprachproduktion*). Diese Ebenen bestehen aus einer Menge von Einheiten oder sogenannten *Knoten*, die die Aktivierung sammeln und an andere Einheiten weitergeben. Bei den meisten Modellen gibt es eine Ebene zwischen der Semantik und der Wortform, wie zum Beispiel eine syntaktische Ebene. Manchmal wird zusätzlich noch eine morphologische Ebene, oder eine Silbenebene (mit Silbenkonstituenten) postuliert. Die Ebene zwischen Semantik und Phonologie kann aber in unterschiedlichen Theorien sehr unterschiedlich aufgebaut sein. Wie oben bereits beschrieben, gibt es die Lemma-Theorie, die besagt, dass auf dieser Ebene die syntaktischen Informationen eines Wortes abgebildet sind, wodurch die orthographischen und phonologischen Informationen auf der nächsten Ebene aktiviert werden. Bonin et al. (1998) beispielsweise meinen, dass sich die konzeptuelle und die semantische Ebene nicht decken, sondern zwei unterschiedliche Ebenen darstellen.

Eine weitere Version ist das Modell von Caramazza (1997) - die *Independent Network theory*, die ich im Folgenden näher darstellen möchte.

4.3.1. Das Modell von Caramazza

Caramazza ändert das oben beschriebene Modell insofern ab, als dass er erstens eine modalitätsneutrale Lemma-Ebene zwischen den semantischen, konzeptuellen Repräsentationen und den Wortformen ablehnt und sich zweitens für eine autonome Ebene ausspricht, die die syntaktische Information der semantischen und Wortformpräsentationen beinhaltet. Er ist der Meinung, dass durch Zugriff auf ein Wort nicht zwingend auch dessen syntaktische Eigenschaften mit abgerufen werden müssen, daher stellen die syntaktischen Eigenschaften eines Wortes eine eigene Ebene dar. Er nennt dieses Modell *Independent Network Model*:

„The Independent Network (IN) model of the lexicon assumes that lexical knowledge is organised in sets of independent networks connected to each other by a modality-specific lexical node.“ (Caramazza 1997, S. 194)

- Das lexikalisch-semantische Netzwerk repräsentiert die Bedeutungen von Wörtern als Set semantischer Eigenschaften.
- Das lexikalisch-syntaktische Netzwerk repräsentiert syntaktische Eigenschaften. Es gibt dabei ein Subnetzwerk mit den Kategorien Verb, Nomen, usw., ein Subnetzwerk mit dem Genus, und weitere Subnetzwerke mit weiteren syntaktischen Merkmalen.
- Da aphasische Patienten manchmal schriftlich Objekte richtig, mündlich aber inkorrekt benennen können, unterscheidet Caramazza ein Netzwerk der Phonemlexeme (mündliche Sprachproduktion) von einem Netzwerk der Graphem-Lexeme (schriftliche Sprachproduktion), die jeweils modalitätsspezifische Repräsentationen enthalten.

Die Semantik greift direkt auf die Wortformnetzwerke zu, und diese Netzwerke, sowie das syntaktische Netzwerk werden gleichzeitig aktiviert. Die Knoten innerhalb eines Subnetzwerks haben auch hemmende Verbindungen, da sie miteinander konkurrieren.

Die Wortproduktion geschieht nun so, dass eine ausgewählte lexikalisch-semantische Repräsentation die lexikalisch-semantische Ebene und die phonologischen und graphematischen Lexem-Netzwerke aktiviert. Die Auswahl eines syntaktischen Knotens ist normalerweise vom semantischen Netzwerk aus nicht ausreichend. Damit aber

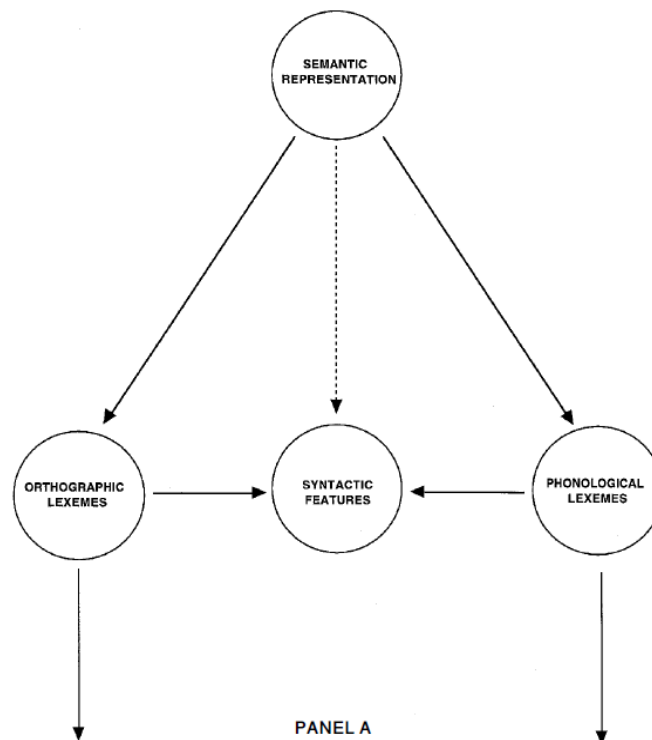


Abb. 15.: Das *Independent Network model* von Caramazza (1997, S. 196). Im Artikel schreibt Caramazza von den *orthographischen* Repräsentationen, in dieser Arbeit wird dies als *graphematische* Repräsentationen bezeichnet.

die Auswahl des ganzen Sets der grammatikalischen Eigenschaften eines Wortes geschieht, muss vorher der modalitätsspezifische lexikalische Knoten (phonologisch oder graphematisch) aktiviert worden sein. Durch dieses Modell versucht Caramazza zu erklären, weshalb manchmal die syntaktischen Merkmale eines Wortes nicht abrufbar sind (durch Hirnschädigungen/Aphasien, oder beim Tip-of-the-tongue-Phänomen) während die Wortform an sich sehr wohl zugänglich ist.

Abbildung 16 stellt eine detaillierte Version des Modells dar, die nur die phonologische Repräsentation umfasst (nicht die graphematische). Der Informationsfluss geht von der höchsten Ebene (Semantik) über die einzelnen Lexeme und deren syntaktische Informationen zu den segmentalen Einheiten.

Somit geschieht hier der lexikalische Zugriff ebenfalls in zwei größeren Schritten. Erstens wird eine modalitätsspezifische, syntaktisch und semantisch spezifizierte Repräsentation ausgewählt, und zweitens wird der phonologische oder graphematische

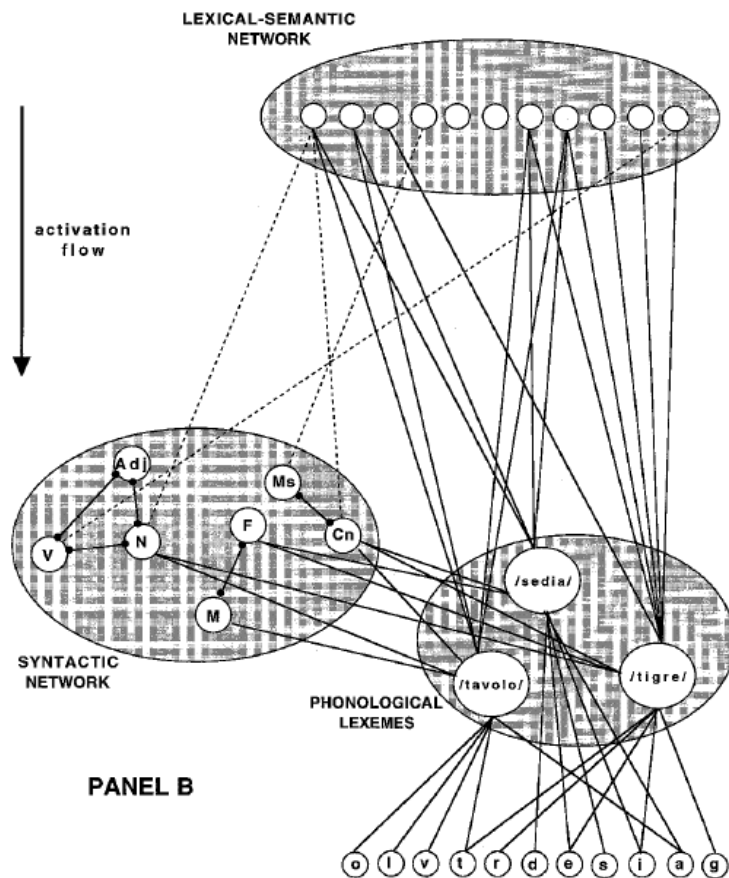


Abb. 16.: Das *Independent Network model* (detaillierter) von Caramazza (1997, S. 197). N=Nomen; V=Verb; Adj=Adjektive; M=maskulin; F=feminin; CN=zählbares Nomen; Ms=Stoffnomen. Gepunktete Linien bedeuten, dass die Aktivierung schwach ist. Verbindungen innerhalb eines Netzwerks sind hemmend.

Inhalt des Lexems ausgewählt (Caramazza 1997). Da im Allgemeinen Bedeutungen mehrere Bestandteile besitzen, werden aufgrund dieser Bestandteile mehrere Lexeme aktiviert (diese Theorie ist allerdings nicht unumstritten). Es wird diejenige Einheit eines Netzwerks weiterverarbeitet, die die höchste Aktivierung erlangt. Leider behandelt Caramazza in diesem Artikel nicht die Frage, ob es in seinem Modell auch ein Feedback gibt, das Informationen der Lexeme zur Syntax oder gar der Semantik bringt.

Modelle, in denen es nur eine Richtung der Verarbeitung gibt (top-down), werden modulare oder serielle Modelle genannt.

4.3.2. Serielle Sprachproduktionsmodelle

Bei einem seriellen Sprachproduktionsmodell werden von der Intention der Aussage bis hin zur Realisierung der Aussage die sprachlichen Elemente Schritt für Schritt nacheinander kodiert und verarbeitet. Wenn der eine Schritt vollzogen wurde, kann zum nächsten Schritt übergegangen werden. Wichtig ist dabei, dass keine Interaktion der einzelnen Stationen vorgesehen ist, sie sind voneinander unabhängig. Ein solches, generelles Modell ist in Abbildung 17 dargestellt.

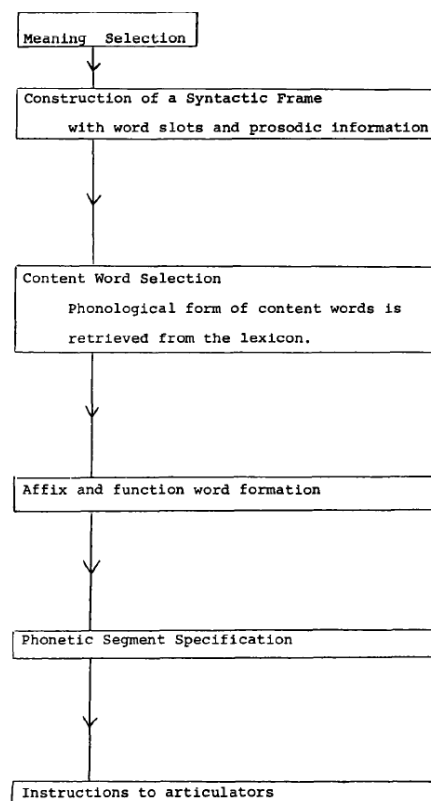


Abb. 17.: Das *Standard Top-Down Serial Processing Model of Speech Production* (Harley 1984, S. 193)

Nach der Ebene der Semantik folgt der syntaktische Rahmen, in dem auch die Betonung kodiert wird. Dann folgen die lexikalische Auswahl, die morphophonemi-

sche Übersetzung und die Planung der Artikulation. In diesem Modell können, wie gesagt, die einzelnen Ebenen nicht miteinander interagieren. Harley (1984, S. 213) beschreibt ein solches Modell als „unbefriedigend“, da es Fehler in der Sprachproduktion oder Versprecher nicht ausreichend beschreiben kann. Lexikalische Alternativen würden hier vor der phonologischen Ebene gelöscht werden, weshalb die phonologische Form der Alternativen den Output nicht mehr beeinflussen sollte, dem ist aber nicht so. Nun wird ein weiteres serielles Sprachproduktionsmodell vorgestellt, das bereits etwas komplexer ist.

4.3.3. Das Sprachproduktionsmodell von Levelt et al. (1999)

Bei dem hier vorgestellten Modell von Levelt handelt es sich um ein serielles Modell zur Beschreibung der Produktion einzelner Wörter. Es umfasst die (semantische) Aktivierung der lexikalischen Konzepte, die Auswahl der Lemmata (was die syntaktische Eigenschaft der Wörter beschreibt), die morphologische und phonologische Kodierung eines Wortes im Kontext der Prosodie und schlussendlich die phonetische Kodierung, die die Artikulation einleitet (Levelt et al. 1999). Die Artikulation selbst wird von dieser Theorie nicht mehr genauer abgedeckt.

Die Autoren entwickelten das Modell hauptsächlich auf Grundlage von Reaktionszeitmessungen. Wie oben bei dem seriellen Modell bereits beschrieben, wird bei diesem Modell die Wortform, d.h. die phonologische Kodierung eines Wortes erst nach Auswahl des Wortes aus dem mentalen Lexikon (mitsamt seiner syntaktischen Merkmale) abgerufen (siehe Abbildung 18).

Die Sprachproduktion läuft auf folgenden Ebenen ab:

- Konzeptuelle Vorbereitung und lexikalisches Konzept (*Conceptual preparation and lexical concepts*):

Alle Wörter mit Bedeutung aktivieren ein ihnen zugrunde liegendes lexikalisches Konzept, dies trifft auf die Wörter der offenen Klasse und auf einige Wörter der geschlossenen Klasse zu. Die Aktivierung dieses Konzepts heißt in dem Modell „konzeptuelle Vorbereitung“. Oft wird das lexikalische Konzept aber im Rahmen einer größeren kommunikativen Absicht aktiviert, wodurch mehrere lexikalische Konzepte aktiviert werden können (Levelt gibt als Beispiel *Stute* an, das eine eigene Bezeichnung für ein weibliches Pferd ist, im Vergleich zu *weiblicher Elefant*, wofür es keine eigene Bezeichnung gibt; 1999,

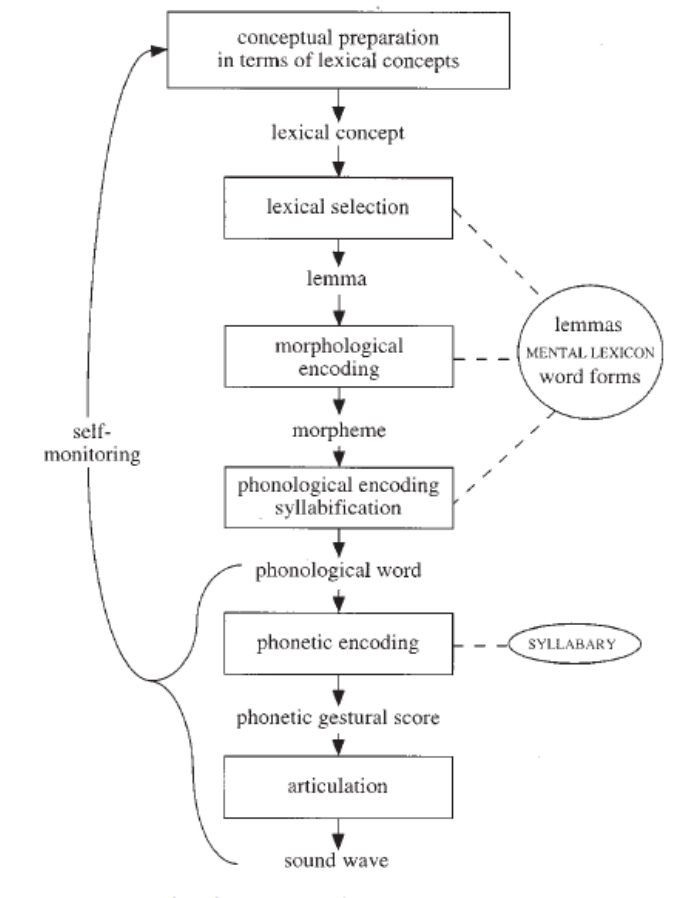


Abb. 18.: Sprachproduktionsmodell (seriell) von Levelt et al. 1999, S. 3

S. 3). Selbst wenn nur ein lexikalisches Konzept aktiviert wird, gibt es verschiedene Möglichkeiten der Wortwahl, da es mehrere Namen für das gleiche Objekt gibt.

- Lexikalische Auswahl und Lemma (*lexical selection and lemma*):
Dies beschreibt den Abruf eines Wortes (bzw. eines Lemmas, das die syntaktischen Eigenschaften dieses Wortes umfasst) aus dem mentalen Lexikon aufgrund des auszudrückenden lexikalischen Konzepts. Im normalen Sprachfluss werden zwei bis drei Wörter pro Sekunde auf diese Art abgerufen, wobei zu beachten ist, dass dabei erstaunlich wenig Fehler gemacht werden. Hier werden auch die syntaktischen Eigenschaften der Wörter abgerufen, sowie Wörter, die keine Bedeutung, sondern nur syntaktische Eigenschaften besitzen, wodurch

die Eingliederung in einen Satz und die Erstellung eines korrekten Satzes selbst möglich wird. Zum Beispiel wird zusammen mit einem Nomen auch gleichzeitig sein Genus erfasst. Ein Verb kann verschiedene Formen annehmen - das lexikalische Konzept wird an die syntaktische Umgebung angepasst (und umgekehrt).

- Morphophonologische Kodierung und Silbentrennung (*morphological encoding, phonological encoding and syllabification*):

Nun wird zur Phonologie übergegangen und die Artikulation vorbereitet. Dabei wird die phonologische Form des Wortes aus dem mentalen Lexikon abgerufen, wodurch manchmal Fehler entstehen können (zum Beispiel das Tip-of-the-tongue-Phänomen - das Wort liegt dem Sprecher „auf der Zunge“, kann aber nicht ausgesprochen werden, die syntaktischen Eigenschaften dieses Wortes sind aber manchmal sehr wohl abrufbar). Der Zugriff auf die Form des Wortes umfasst drei verschiedene Bereiche: die Morphologie, die Prosodie und die segmentale Ausstattung. Das Wort kann in seine Morpheme aufgeteilt werden, die prosodischen Informationen, die Anzahl der Silben und das Betonungsmuster der Wortteile werden abgerufen. Die Wörter werden in einzelne Laute aufgeteilt. Levelt und Kollegen (1999) weisen darauf hin, dass die Wörter nicht zuerst in Silben und dann in Laute aufgetrennt werden, da die Silbentrennung nicht im mentalen Lexikon abgespeichert ist. Die Phoneme eines Morphems werden gleichzeitig bereitgestellt (aber in korrekter Reihenfolge). Bei der Silbentrennung handelt es sich um einen relativ späten Prozess, da er sehr stark von der sich verändernden Umgebung eines Wortes abhängt (und, wie in Kapitel 2.2.3 *Silbengrenzen* bereits beschrieben, fallen die Morphemgrenzen nicht zwangsläufig mit den Silbengrenzen zusammen). Dabei werden die in jenem Kapitel beschriebenen Regeln wie die des maximalen Onsets eingehalten.

- Phonetische Kodierung (*phonetic encoding*):

Die Theorie umfasst diesen Bereich nicht mehr ganz, sondern beschreibt nur die Berechnung der artikulatorischen Gesten eines phonologischen Wortes auf verschiedenen Ebenen: glottal, nasal und oral. In diesen Bereich fällt auch das Silbenlexikon. Levelt und Kollegen (1999) nehmen an, dass der Sprecher Zugriff auf die artikulatorischen Gesten hochfrequenter Silben hat. Dafür spricht, dass die Koartikulation eher innerhalb einer Silbe als zwischen Silben auftritt. Da

beim Großteil des Gesprochenen nur eine verhältnismäßig geringe Silbenanzahl verwendet wird (nach Levelt und Kollegen im Englischen nicht mehr als ein paar Hundert Silben; 1999, S. 5), scheint es sinnvoll, dass dem Sprecher die häufigsten Silben direkt zur Verfügung stehen, wodurch sie nicht immer wieder aus einzelnen Lauten neu erstellt werden müssen - deswegen das Silbenlexikon. Im Modell von Levelt und Kollegen (1999) wird das Silbenlexikon direkt durch die phonologischen Segmente angesprochen.

- Die Selbstüberwachung oder -steuerung (*Self-Monitoring*): Dies stellt einen weiteren wichtigen Faktor bei der Sprachproduktion dar. Es geht dabei darum, dass wir dauernd selbst kontrollieren, was wir sagen und hören, und wir überprüfen unsere Produktionen auf Fehler, Unflüssigkeiten etc. Levelt und Kollegen schreiben, dass wir auch unsere „interne Sprache“ in diesem Sinne überprüfen.

Dieses Modell beschreiben Levelt und Kollegen als ein *feedforward activation-spreading network* (1999, S. 6), was so viel bedeutet wie ein „die Aktivierung weiterleitendes Netzwerk“ (siehe Abbildung 19 - veranschaulicht am englischen Wort *escorting*).

In diesem Netzwerk gibt es drei Schichten. Die erste Schicht ist die konzeptuelle Schicht, die auch die lexikalischen Konzepte umfasst. Jedes lexikalische Konzept stellt einen eigenen Knoten dar und kann mit anderen Knoten verbunden sein. Zusätzlich ist das lexikalische Konzept mit dem syntaktischen Lemma verbunden, das sich bereits in der nächsten Schicht befindet. Hier werden die syntaktischen und grammatischen Eigenschaften des Wortes kodiert (Numerus, Tempus etc.). „Each word in the mental lexicon, simple or complex, content word or function word, is represented by a lemma node“ (Levelt et al. 1999, S. 6). Nachdem das Lemma ausgewählt wurde, wird die dritte Schicht aktiviert, die die Wortform betrifft. Hier gibt es Verbindungen zu Morphem-Knoten, die wiederum mit ihren dazugehörigen (nummerierten) Segment-Knoten verbunden sind. Ein lexikalischer Eintrag umfasst somit laut Levelt et al. (1999) ein Item im mentalen Lexikon, das aus einem Lemma, einem lexikalischen Konzept und seinen Morphemen mitsamt segmentaler und metrischer Eigenschaften besteht.

Wie wird nun ein Wort ausgewählt? Der Abruf geschieht durch Erhöhung des Aktivierungslevels des lexikalischen Zielkonzepts, das einen Teil der Aktivierung an

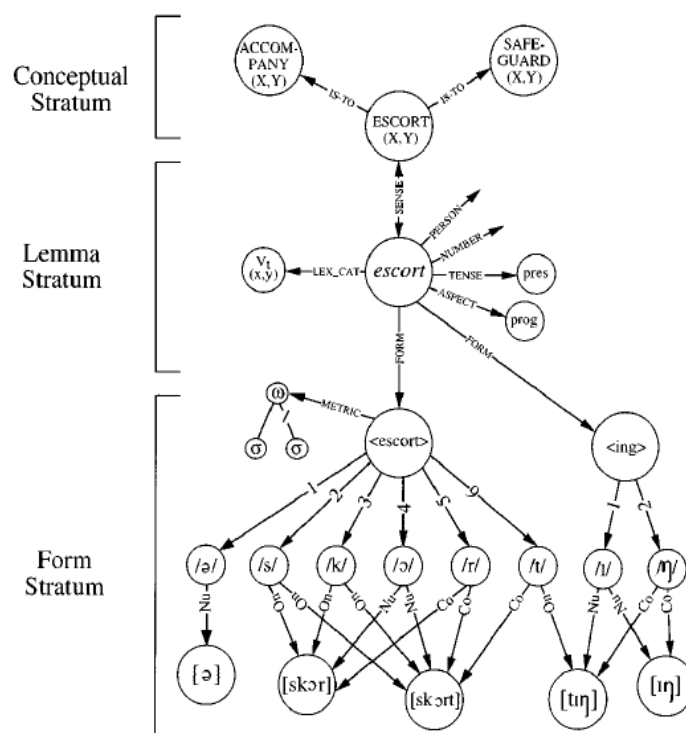


Abb. 19.: Fragment des lexikalischen Netzwerks beim Zugriff aufs Lexikon (Levelt et al. 1999, S. 4)

benachbarte Knoten in der gleichen Schicht, sowie an Knoten in der nächsten Schicht (Lemma-Schicht) weitergibt. Der Lemmaknoten mit der höchsten Aktivierung wird weiterverarbeitet. Dies geschieht auch bei Funktionswörtern ohne semantischen Inhalt.

Im Vergleich dazu wollen wir nun auf andere Sprachproduktionsmodelle eingehen, die eine Interaktion zwischen den einzelnen Ebenen zulassen.

4.3.4. Interaktive Sprachproduktionsmodelle

Die Erstellung von Sprachproduktionsmodellen basiert zum großen Teil auf beobachteten Fehlern in der Sprachproduktion. Verschiedene Faktoren können zu Fehlern und Versprechern führen. Einige Autoren meinen, dass auch die Phonologie der Wörter oder Wörter, die der Sprecher gerade liest, während er spricht, seine Äußerung beeinflussen. Aus diesem Grund müsste man andere Modelle entwerfen, die auch jene

Fehler erklären konnten (mehr dazu unter Dell 1986, Dell 1966, Harley 1984).

„It is found that phonological similarity between the target and intrusion is a motor determinant of error occurrence.“ (Harley 1984, S. 191)

Fehler in der Sprachproduktion können auf phonologische Faktoren zurückgeführt werden. Auch gibt es Fehler, die eine gemischte phonologische und semantische Grundlage haben können. Diese kombinierten Fehler kommen häufiger vor, als sie durch reinen Zufall auftreten würden. Dabei ist unklar, ob es sich um phonologisch bedingte semantische Fehler, oder um semantisch bedingte phonologische Fehler handelt. Entscheidend ist jedoch, dass sie zum gleichen Fehlertypus beitragen (Harley 1984, S. 210f).

In einem interaktiven oder konnektionistischen Modell geht man von einem zentralen Lexikon aus und es besitzt drei Ebenen (wie oben beschrieben): Semantik, Lexeme und Phoneme. Zwischen den einzelnen Ebenen gibt es mehrere Verbindungen, und bis ein Wort tatsächlich ausgesprochen wird, kann die Aktivierung der einzelnen Knoten mehrmals hin- und hergegangen sein (Huber et al. 2006).

In Kapitel 3.4 *Reime bei aphasischen Erwachsenen* wurden bereits ein paar Studien vorgestellt, die ein solches interaktives Sprachproduktionsmodell postulierten (Spencer et al. 2000, Blumstein et al. 2000). Nun soll ein solches interaktives Sprachproduktionsmodell genauer erläutert werden.

4.3.5. Das Sprachproduktionsmodell von Dell (1986)

Im Gegensatz zu dem oben beschriebenen seriellen Modell von Levelt et al. (1999) handelt es sich bei dem Modell von Dell (1986) um ein nicht-serielles Modell, das eher als *Aktivationsflussmodell* (Levelt et al. 1999) beschrieben werden kann (*spreading-activation model/connectionist model/interactive activation model*):

„The theory combines assumptions from linguistic theory regarding the units and structures of language with a precisely specified retrieval mechanism based on spreading activation.“ (Dell 1986, S. 283)

Dabei wird ein „Netzwerk von linguistischen Regeln und Einheiten [angenommen], in denen Entscheidungen darüber, welche Einheiten oder Regeln gewählt werden sollen, auf den Aktivierungsregeln der Knoten basieren, die diese Regeln oder Einheiten repräsentieren“ (Dell 1986, S. 283).

Dell beschreibt das Lexikon als Netzwerk und unterteilt es in die Ebene der Semantik, der Syntax, der Morphologie und der Phonologie (in dieser Reihenfolge). Dieses Netzwerk hat je nach Ebene Knoten für semantische Konzepte, Wörter, Morpheme, Phoneme, phonologische Merkmale. Die konzeptuellen Knoten in der Semantik sind mit Wortknoten der Syntaxebene verbunden, die Wortknoten mit den Morphemknoten, die Morphemknoten mit Phonemknoten, die Phonemknoten mit Merkmalen (velar, labial usw.). Die Ebenen können unterschiedlich groß sein und jede Ebene besitzt eigene generative Regeln. Diese Regeln bestimmen die Kombinationsmöglichkeiten der Einheiten dieser Ebene (und somit auch deren Produktivität).

Die Produktion von Sprache, wie beispielsweise in Form eines Satzes, ist auf jeder Ebene gleichzeitig repräsentiert, aber nicht im gleichen Ausmaß. Der Satz ist jeweils als geordnetes Set von Wörtern, Morphemen bzw. Phonemen repräsentiert, die durch Knoten dargestellt werden. Die Konstruktion dieser Repräsentation läuft auf allen Ebenen ab, aber „the selection of items for a lower representation must await the construction of corresponding structures of the higher representation.“ (Dell 1986, S. 287). Die Knoten einer höheren Ebene aktivieren die Knoten der direkt darunterliegenden Ebene. Das bedeutet, dass erst nach Auswahl eines Wortes auf der syntaktischen Ebene dieses Wort phonologisch kodiert werden kann. Dabei stellt jede Ebene einen lexikalischen Auswahlprozess dar.

Ganz oben in der Hierarchie steht hier die Semantik. Sie steuert, was gesagt werden soll und beinhaltet die gedanklichen Konzepte, die zu geeigneten Zeitpunkten während der syntaktischen Kodierung zugänglich gemacht werden (Dell 1986, S. 287). Die Semantik wird in der Theorie nicht weiter behandelt. Die unteren Ebenen der Syntax, Morphologie und Phonologie bilden jeweils einen Rahmen bzw. ein geordnetes Set von Kategorien, und mit bestimmten Regeln werden die Plätze in diesem Rahmen gefüllt - je nach Aktivierungsgrad der Knoten der Items, die eingefügt werden sollen. Es wird das Item genommen, dessen Knoten den höchsten Aktivierungsgrad besitzt. Nach Auswahl des Items wird sein Aktivierungslevel auf Null gesetzt. Wenn zum Ausdruck eines bestimmten semantischen Konzepts mehrere Realisierungen möglich sind (zum Beispiel die Wörter *sinken* und *ertrinken*), dann konkurrieren diese Varianten miteinander. Es kommt das Item in den Position des Rahmens, das die höchste Aktivierung erlangt.

Wenn der Aktivierungsgrad einer Einheit größer als Null ist, schickt dieser Knoten einen Anteil seiner Aktivierung an alle Knoten, mit denen er verbunden ist (*sprea-*

ding). Dieser Anteil muss nicht für jede Verbindung gleich groß sein. Bei Ankunft dieser Aktivierung beim Zielknoten wird dies zu dem dortigen Aktivierungslevel hinzugefügt (*summation*). Im Laufe der Zeit sinkt der Aktivierungsgrad (*decay*). *Spreading* (Verteilung), *summation* (Summierung), und *decay* (Abfall/Abbau) kommt zu jeder Zeit auf allen Knoten im lexikalischen Netzwerk vor, egal ob der Knoten ein Teil der Repräsentation oder der semantischen Intention ist, oder nicht. So kann es vorkommen, dass bereits kodierte Phoneme eines Wortes noch einmal verarbeitet werden, da andere aktivierte Knoten auch noch mit diesen Phonemen verbunden sind (Dell 1988).

Im Folgenden ist eine Graphik zu sehen, die einen Zeitpunkt in der Produktion des Satzes „Some swimmers sink“ darstellt (Abb. 20).

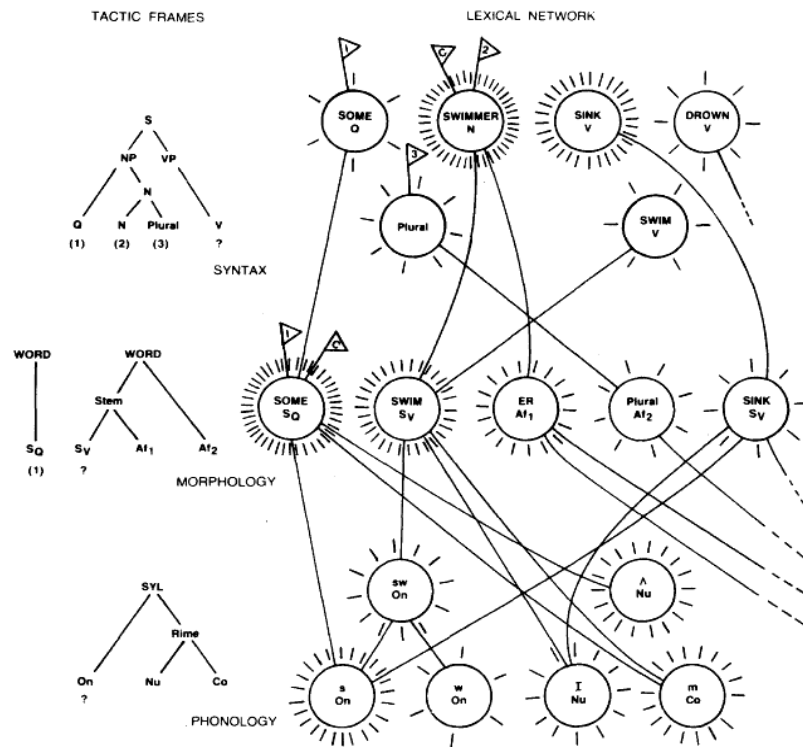


Abb. 20.: Lexikalisches Netzwerk bei der Sprachproduktion (Dell 1986, S. 290)

Die Modelle auf der linken Seite der Abb. 20 (die sogenannten Rahmen - *tactic frames*) stellen auf jeder Ebene ein geordnetes Set von Plätzen dar, die bereits gefüllt wurden. Auf jeden Knoten im lexikalischen Netzwerk, der für ein Item im Rahmen

steht, wurde eine Fahne gesetzt. Ein Fragezeichen steht für die Position in jedem Rahmen, der momentan gefüllt wird. Die Markierungen um die Knoten herum stellen den Aktivierungsgrad dar und eine Fahne mit „C“ markiert den momentanen Knoten einer jeden Ebene. Die Knoten wurden bezüglich ihrer Zugehörigkeit zu einer Kategorie markiert: Syntaktische Kategorien für Wörter sind (Q)uantifizierer, (N)omen, (V)erb, Pluralmarkierung; morphologische Kategorien für Morpheme sind (S)tamm, (Af)fix, und phonologische Kategorien für Laute sind (On)set, (Nu)kleus, (Co)da. Viele Knoten wurden der Einfachheit halber ausgelassen (Silben, Silbenkonstituenten, phonetische Merkmale).

Wichtig für die Untersuchung im Rahmen dieser Arbeit ist die Tatsache, dass in dieser Theorie die Verbindungen der Knoten nicht nur in eine Richtung gehen, sondern auch wieder zurück. „Because of the bottom-up as well as the top-down connections in the network, lexical retrieval is highly interactive“ (Dell 1986, S. 317). Es gibt ein Bottom-up-Feedback, das Entscheidungen in höher liegenden Ebenen beeinflusst. Dell (1986, S. 317) gibt dazu folgendes Beispiel: *I'm writing a letter to my mother*. Es wird angenommen, dass nach Kodierung von *I'm writing a* nun das Wort *letter* aktiviert wird - dabei werden aber auch andere Repräsentationen auf anderen Ebenen aktiviert, wie zum Beispiel:

- semantische oder pragmatische Relationen (Wörter wie *note*, *stamp*, *envelope*),
- Nomen anderer Teile der Äußerung (das Wort *mother*),
- Nomen der nichtsprachlichen Umgebung (zum Beispiel das Wort *dog*, falls gerade ein Hund das Zimmer betritt),
- Nomen, die phonologisch ähnlich zu *letter* sind, wie *lettuce*, *litter*.

Diese Wörter konkurrieren nun miteinander. Es kann sein, dass nicht *letter* realisiert wird, sondern ein anderes aktiviertes Wort. Es kann auch geschehen, dass zwei Wörter in etwa gleich aktiv sind und dadurch ein Mischwort realisiert wird, wodurch zwei Items für die gleiche Position im Rahmen ausgewählt werden. Dell (1986, S. 318) schreibt weiter, dass solche Vermischungen (wie zum Beispiel *Let's stop* statt *Let's start*) ein semantischer UND phonologischer Fehler sind - die Aktivierungsquellen wurden kombiniert. Versprecher, Fehler in der Sprachproduktion können gleichzeitig mehrere Ursachen haben. Dies stellt natürlich einen Gegensatz zu seriellen Modellen der Sprachproduktion dar.

In einem weiteren Artikel, der eine Weiterentwicklung des oben beschriebenen Modells vorstellt, präsentiert Dell (1988) die Sprachproduktion als zweiteiliges Netzwerk (ebenso wie die meisten Autoren, siehe die Einleitung zum Kapitel 4.3 *Sprachproduktionsmodelle*). Es handelt sich dabei um ein lexikalisches Netzwerk und um ein Wortformnetzwerk. Dabei wird die phonologische Struktur eines Wortes auf zwei Komponenten aufgeteilt:

- auf einen Rahmen, der eine Sequenz von Positionen darstellt, die die abstrakte Form des Wortes repräsentieren - die Anzahl und Art der Silben und Phoneme des Wortes (CV, CVC),
- und auf eine eigene Repräsentation der Laute des Wortes, die mit diesen Positionen verbunden sind.

Zur Veranschaulichung siehe Abbildung 21 bei der Produktion von *deal back*. Der Einfachheit halber werden in dieser Abbildung nur die Wort- und Phonemebene dargestellt. Das aktuelle Wort ist *deal*, alle Verbindungen zwischen den Knoten sind aktivierend und gehen in beide Richtungen (positives Feedback). Verbindungen zwischen dem lexikalischen und den Wortformnetzwerken werden durch punktierte Linien angezeigt. Eine wichtige Änderung im Vergleich zum ersten Modell stellt der *header* dar: Er repräsentiert das Muster der Phonemkategorien des Wortes (CVC, CV). Die Komponente C ist dann wieder mit den einzelnen Phonemen verbunden. Der am meisten aktivierte Header wählt die Phoneme für jede seiner Kategorien aus. In dem Beispiel von Dell wäre der CVC-Header der aktivierteste: Er wählt das CVC-Muster, und somit den Konsonanten, einen Vokal und wieder einen Konsonanten mit der höchsten Aktivierung.

„The spread of activation allows for activated words to retrieve their constituent phonemes in the lexical network“ (Dell 1988, S. 129). Die Aktivierung findet allerdings in beide Richtungen statt, d.h. auch andere Wörter und Phoneme werden aktiviert und tragen zum Abrufprozess bei.

Einige Jahre später stellten Dell und O’Seaghdha (1997) ein Modell vor, das eine Kombination von seriellen und interaktiven Modellen darstellt. Hier ist wieder die Einteilung in zwei Schritten zu erkennen (Semantik, Phonologie). Dieses Modell soll im Allgemeinen modular (d.h. seriell), stellenweise aber interaktiv fungieren. Es gibt keine hemmenden Verbindungen, Top-down und Bottom-up-Verbindungen sind im

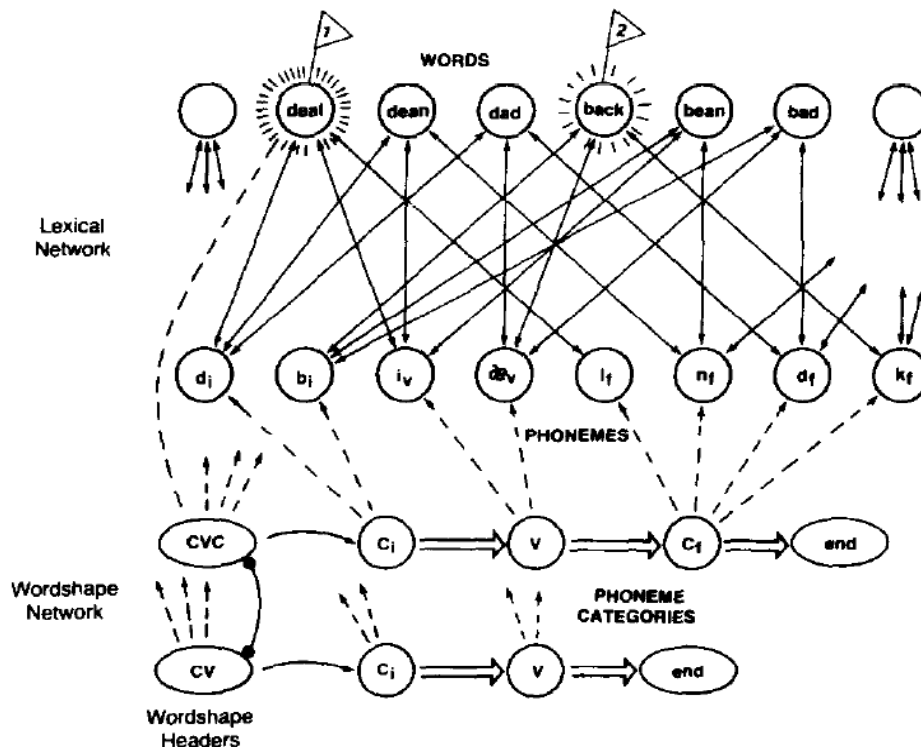


Abb. 21.: Das lexikalische und das Wortformnetzwerk (Dell 1988, S. 128)

ganzen System möglich. Die Phonologie hat aber einen geringeren Effekt auf die Semantik als umgekehrt. Unter Abbildung 22 ist ein Modell zu sehen, bei dem die gemeinsamen semantischen Eigenschaften von *cat*, *dog* und *rat* grau hinterlegt sind.

4.3.6. Die schriftliche Sprachproduktion

In den oben genannten Modellen ist meist nur von der wörtlichen Sprachproduktion oder von allgemeiner Sprachproduktion die Rede. Auf die schriftliche Sprachproduktion wird nur sehr wenig eingegangen, wie auch von Rapp und Goldrick (2000) bemerkt wird. Auch Miceli et al. (1997, S. 37) sind der Meinung, dass „[...] there is almost no research on the mechanisms of orthographic word production [...]“. Bonin et al. (1998, S. 269) betonen: „Unlike speech production, lexical access in written production has not systematically been investigated experimentally“. Aus diesem Grund - und da das in dieser Arbeit durchgeführte Experiment schriftlicher Modalität ist - möchte ich gerne die schriftliche Sprachproduktion etwas näher behandeln.

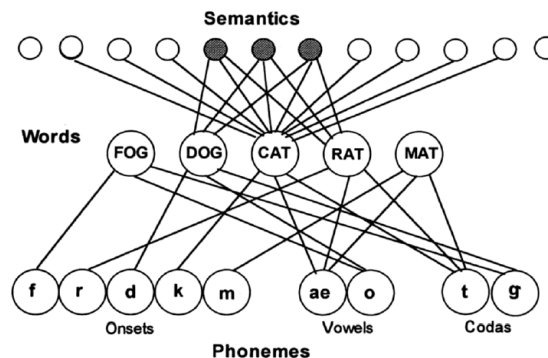


Abb. 22.: Lexikalisches Netzwerk bei der Sprachproduktion (Dell und O'Seaghdha 1997, S. 805)

Zuallererst ist eine wichtige Frage zu klären: Wie hängen mündliche und schriftliche Sprachproduktion zusammen? Da sowohl der Mensch als Gattung, als auch jeder Mensch im Heranwachsen zuerst sprechen und dann schreiben lernt(e) (Rapp et al. 1997), könnte man annehmen, dass die schriftliche Sprachproduktion von der mündlichen abhängt. Werden also beim Schreiben zuerst immer die phonologischen Wortformen kodiert, bevor die graphematischen Wortformen verarbeitet werden können, oder stellen die mündlichen und schriftlichen Repräsentationen zwei unterschiedliche Output-Ebenen dar? Die beiden Versionen sind in Abbildung 23 abgebildet.

Caramazza (1997) spricht sich für die direkte Beziehung der Lemmaknoten zu den orthographischen Wortformen aus und ist entgegen der Annahme, dass zuerst eine Verarbeitung über die phonologischen Wortformen erfolgen muss. Man nimmt an, dass der Zugriff auf die orthographischen Wortformen unabhängig von dem Zugriff auf die phonologischen stattfindet, da es Patienten gibt, die in einer Modalität bessere Leistungen erzielen, als in der anderen: Sie haben beispielsweise Probleme beim Sprechen, können aber Schreiben, wobei in den Fällen post-phonologische Bereiche wie die Artikulation als Störungsbereiche ausgeschlossen werden können (Caramazza 1997, S. 189ff; Miceli et al. 1997; Shelton und Weinrich 1997; Rapp et al. 1997; siehe dazu auch das Logogenmodell in Kapitel 4.1 *Das mentale Lexikon*). Die Verbindung vom Lemma zu einem phonologischen oder einem orthographischen Lexem muss ebenfalls nicht identisch sein.

Rapp et al. (1997) meinen aber, dass der Umstand, dass schriftliche Sprachpro-

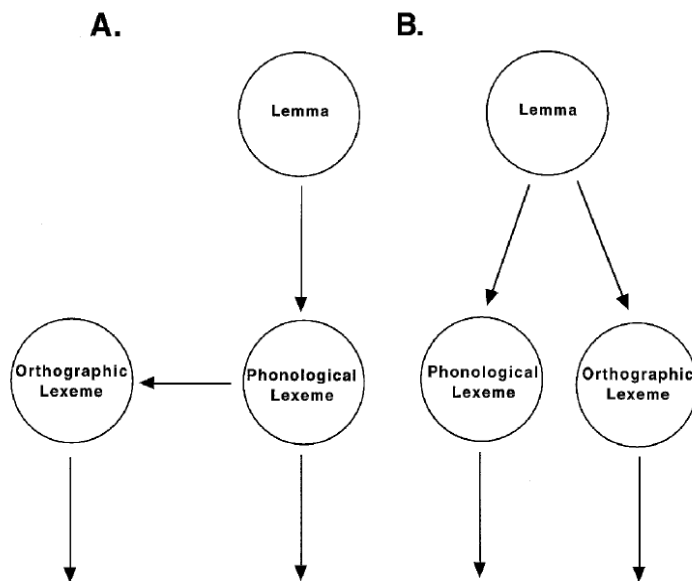


Abb. 23.: Darstellung der phonologischen und orthographischen Repräsentationen (Caramazza 1997, S. 190)

duktion nicht zwingend über die phonologischen Repräsentationen gehen *muss*, nicht sofort bedeutet, dass die Phonologie als Übermittler nicht benutzt werden *kann*. Als Evidenz, die eine optionale Verwendung der Phonologie beim Schreiben befürwortet, betrachten Rapp et al. (1997, S. 99), dass viele Sprecher die subjektive Erfahrung machen, sich beim Schreiben „selbst zu hören“. Selbst unbeeinträchtigte Sprecher machen des öfteren Fehler beim Schreiben homophoner Wörter mit unterschiedlicher Schriftweise: Engl. *here* und *hear* (Rapp et al. 1997, S. 100) - diese Fehler werden in der Literatur als *slips of the pen* bezeichnet (Bonin et al. 1998). Dies kann entweder als einfache Begleiterscheinung (es werden schriftliche und mündliche Wortformen produziert), oder als phonologisch vermitteltes Schreiben (der Prozess des Schreibens verläuft über Repräsentationen der Phonologie) interpretiert werden. Bonin et al. (1998) schreiben, dass es im Französischen, wie sowohl auch im Englischen, viele stumme Buchstaben gibt, und die Leute können dennoch richtig schreiben².

²Hierbei ist allerdings die Regelmäßigkeit im Schriftsystem im Vergleich zu Ausnahmen zu betrachten - zum Beispiel im Deutschen sind zwei Konsonanten nach einem kurzen Vokal regelhaft (*Messer*), während ein einzelner Konsonant nach einem langen Vokal auftritt (oder umgekehrt, in: *Meter*). Solche Regeln gehören zum Schriftspracherwerb dazu und stellen keine Ausnahmen dar. Eine Ausnahme wäre der kurze Vokal bzw. der einzelne Konsonant in *Bus* - diese Wortform

Auch Bonin und Kollegen (1998, S. 283) haben durch ihre Experimente herausgefunden, dass „[...] phonological information is not necessarily needed for the retrieval of orthographic information“.

Miceli et al. (1997) stellen zwei unterschiedliche Hypothesen vor: eine starke und eine schwache Version der graphematischen Autonomiehypothese. Bei beiden Hypothesen werden die graphematischen Lexeme direkt von der Semantik ausgewählt.

- Die schwache Version besagt, dass „orthographic and phonological lexemes are directly and independently activated by lexical semantic representations, and that phonological and orthographic lexemes are directly connected to each other [...] the activation of one lexical form — either phonological or orthographic — will lead to the activation of its corresponding lexeme in the other modality“ (S. 65f). Es gibt somit von der Semantik ausgehend eine direkte Aktivierung zur gewünschten Modalität und eine indirekte zur anderen, dadurch wird das gleiche Wort phonologisch und graphematisch aktiviert.
- Die starke Version stimmt nur im ersten Teil mit der schwachen Version überein. Es gibt keine direkten Verbindungen zwischen phonologischen und orthographischen Formen, „phonological and orthographic lexical forms receive input both from the semantic component and from sublexical orthography–phonology and phonology–orthography conversion procedures, respectively“ (S. 66): Die Interaktion zwischen phonologischem und graphematischem Lexikon erfolgt somit nur über Konversionsrouten, die eine Art Zwischenschritt von phonologischen zu graphematischen Formen und umgekehrt darstellen.

Im Allgemeinen werden bei der schriftlichen Sprachproduktion zwei verschiedene Verarbeitungswege unterschieden (vgl. Bonin et al. 1998; Houghton und Zorzi 2003; Kreiner 1992):

- Die eine Route führt über das Lexikon, in dem die Wortformen bereits gelernter Wörter repräsentiert sind und somit wortspezifische graphematische Informationen abgerufen werden können. Da es sich um das Lexikon handelt, ist die Frequenz der Wörter für ihren Abruf von Bedeutung.

muss extra im Lexikon abgespeichert werden.

- Die andere Route ist die Phonem-Graphem-Konversion. Es handelt sich hierbei um eine Regelstrategie, um das Wissen, als welche Buchstaben und Buchstabenkombinationen die Phoneme oder Laute einer Sprache repräsentiert sind.

Ein Beispiel zur zweiten Route gibt Kreiner (1992): Im Englischen mappt der Laut /k/ auf die Buchstaben K, C, und CH. Wenn ein Sprecher (und Rechtschreiber) des Englischen nun ein ihm unbekanntes, gehörtes Wort niederschreiben soll in dem der Laut /k/ vorkommt, wird er sich auf die Wahrscheinlich berufen, wie oft welches Graphem (bzw. welcher Buchstabe) welchen Laut repräsentiert - und diesen Buchstaben dann auswählen und niederschreiben. Die Regelstrategie verbraucht wenig Speicherplatz, führt aber zu vielen Fehlern, da die Sprache häufig inkonsistent ist. Sie wird bei unbekannten und nicht gut gelernten Wörtern angewendet, Ausnahmen werden hier regularisiert und führen zu Fehlern (Houghton und Zorzi 2003).

Die lexikalische Strategie kann man sich wie ein Wörterbuch vorstellen. Die gesamte Schreibweise eines Wortes ist im mentalen Lexikon repräsentiert, die einzelnen Buchstaben sind in der richtigen Reihenfolge angeordnet. Darin befinden sich alle Wörter, die die jeweilige Person kennt oder gelernt hat, sowie deren richtige Schreibweise (dies ist besonders wichtig bei Ausnahmen, da diese ja über die reguläre Route nicht korrekt geschrieben werden können). Seltene Wörter sind dabei schwieriger zu schreiben als häufigere. Im Lexikon befindet sich keine Information darüber, wie die Grapheme nun als Buchstaben niedergeschrieben werden sollen - es handelt sich eher um abstrakte graphematische Einheiten (und ihre Beziehung zu den Buchstaben kann wie die Beziehung der Phoneme zu den Lauten verstanden werden; Kreiner 1992).

Kreiner (1992) meint (und bezieht sich in seinem Artikel auf das Englische), dass beide Routen wichtig für korrektes Schreiben sind und parallel benutzt werden. Eine Route alleine könnte korrektes Schreiben nicht erklären. Beide Strategien können auch zu verschiedenen Teilen desselben Wortes beitragen.

Auch Houghton und Zorzi (2003) stellen ein Modell mit zwei Routen vor (*dual-route model of spelling*), das sich in großen Teilen mit dem von Kreiner beschriebenen Modell deckt. Houghton und Zorzi bezeichnen die Phonem-Graphem-Konversion als *direkte Route*: Es handelt sich dabei um ein zweischichtiges, assoziatives Netzwerk, das direkte Verbindungen von silbisch strukturierten phonologischen (Input-) zu graphematischen (Output-) Repräsentationen besitzt. Die zweite (lexikalische) Route ist

frequenzsensitiv: Knoten frequenterer Wörter werden schnell aktiviert. Es wird parallel auf beide Routen zugegriffen und es gibt Interaktionen zwischen ihnen. Das, was schlussendlich niedergeschrieben wird, ist ein kombinierter Output beider Routen. Das Zweiroutenmodell für die schriftliche Sprachproduktion ist in Abbildung 24 dargestellt.

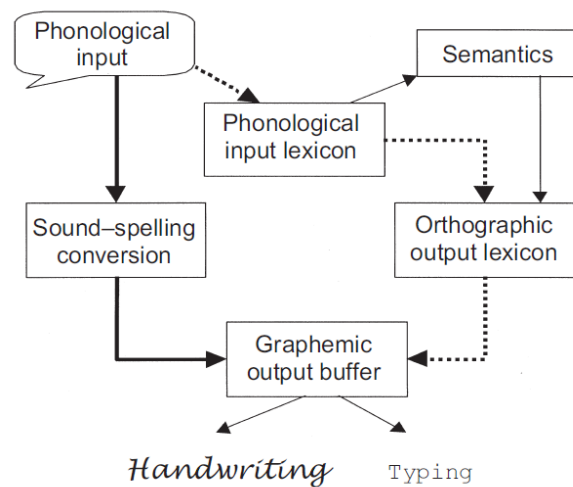


Abb. 24.: Das *dual-route model for spelling* von Houghton und Zorzi (2003, S. 117). Punktirierte Pfeile repräsentieren die lexikalische Route, fett markierte Pfeile die Phonem-Graphem-Konversion. Es wird auch ein direkter Input von der Semantik dargestellt.

In diesem Bereich wird sicher noch viel weitere Forschung nötig sein und ich verweise hiermit auf weitere Literatur, da in diesem Kapitel nur ein genereller Überblick über die relevanten Faktoren bei der schriftlichen Sprachproduktion gegeben werden sollte: Näheres über das Zweiroutenmodell oder Fehler beim Schreibprozess siehe beispielsweise Glasspool und Houghton (2005), mehr zu einem interaktiven Modell bei Houghton und Zorzi (2003); ein weiterer wichtiger Forschungsschwerpunkt ist die Frage, ob tatsächlich ein graphematisches Input-Lexikon und ein graphematisches Output-Lexikon existieren, oder ob es ein Lexikon für Schriftsprache allgemein gibt. Näheres dazu finden Sie unter Hyesuk (2011) und Ring (2001).

4.4. Zusammenfassung

In diesem Kapitel wurde das mentale Lexikon, der Zugriff darauf und in diesem Sinne Sprachproduktionsmodelle vorgestellt, die die mündliche und die schriftliche Produktion näher erläuterten. Anhand der vorgestellten Studien und Theorien kann man sagen, dass ein interaktives Sprachproduktionsmodell den Sprachproduktionsprozess realitätsgetreuer darzustellen scheint als ein serielles Modell. Auch wurde darüber diskutiert, inwiefern Phonologie und Schriftsprache zusammenhängen.

5. Experiment

In diesem Kapitel möchte ich nach einem kurzen Einblick in Priming-Experimente das Experiment von Rapp und Samuel (2002) vorstellen, das ich für das Deutsche adaptiert und mit deutschen L1-Sprechern durchgeführt habe. Die Erstellung der Items (siehe Anhang) und der Verlauf des Experiments werden dokumentiert, sowie die Ergebnisse am Ende vorgestellt.

5.1. Priming-Experimente

„Priming is one of the most basic expressions of human memory, influencing how we perceive and interpret the world.“ (Henson 2003, S. 54)

Lupker und Williams (1989) beschreiben Priming-Experimente folgendermaßen:

„In a typical priming task, two stimuli are presented sequentially. The first, or prime, stimulus establishes a particular context. Empirically, the question is whether this context affects processing of the second, or target, stimulus.“ (Lupker und Williams 1989, S. 1033)

Beim Priming handelt es sich im Rahmen sprachwissenschaftlicher Forschung um eine Technik, in der einer Person ein bestimmter Stimulus präsentiert wird (visuell oder auditiv), wodurch sich bei darauffolgender Präsentation des gleichen oder eines weiteren Stimulus' die Verarbeitung dieses zweiten Stimulus' bei dieser Person bezüglich Geschwindigkeit, Tendenz oder Richtigkeit verändert („[...] change in the speed, bias or accuracy of the processing of a stimulus [...]“ in: Henson 2003, S. 54; Schacter und Buckner 1998):

- Über den Stimulus kann schneller entschieden werden (z.B. bei lexikalischem Entscheiden),
- bei Produktionsaufgaben steigt die Tendenz, ihn zu produzieren,

- oder bei Identifikationsaufgaben wird er besser bzw. häufiger korrekt erkannt (z.B. bei nur kurzer oder teilweiser Präsentation des Stimulus).

Wenn dies der Fall ist und eine Veränderung der Reaktion des Probanden eintritt, spricht man von einem *Priming-Effekt*, der sozusagen eine Erleichterung des Verarbeitens bedeutet (Weldon 2000). Manchmal wird auch einfach der Begriff „Priming“ für „Priming-Effekt“ verwendet.

Der erste Stimulus wird dabei *Prime* genannt. Der zweite, nachfolgende Stimulus heißt entweder *geprimter Stimulus* (oder auch *Target - Ziel*), wenn sich davor ein Prime befindet - oder *nicht geprimter Stimulus*, wenn sich davor kein Prime befindet. Priming-Experimente können relativ einfach aufgebaut sein: Es wird lediglich die Verarbeitung eines geprimten Stimulus' mit der eines nicht geprimten Stimulus' verglichen (Henson 2003).

Bei Priming-Experimenten können sowohl Reaktionszeitmessungen durchgeführt werden, als auch quantitative und qualitative Analysen der Antworten des Probanden. Dabei wird eher die unbewusste Verarbeitung oder Erinnerung angesprochen, als die bewusste. Man spricht in diesem Zusammenhang von dem „impliziten“ oder „nicht-deklarativem Gedächtnis“: Es geht um „[...] a nonconscious influence of past experience on current performance or behavior“ (Schacter und Buckner 1998, S. 185). Beispielsweise amnestische Personen konnten sich nicht bewusst an den Prime erinnern, zeigten aber einen Priming-Effekt wie gesunde Probanden (Warrington und Weiskrantz 1974, in: Henson 2003).

Beim Priming gibt es verschiedene Typen, die unterschieden und miteinander kombiniert werden können. Ausschlaggebend ist dabei die Beziehung vom Prime zum Stimulus (Gordon und Baum 1994, S. 662; Henson 2003, S. 55ff; Neill 1997, Schacter und Buckner 1998):

- Phonologisches Priming: Der geprimte Stimulus steht in einem Formzusammenhang mit dem Prime.
- Semantisches Priming: Der geprimte Stimulus steht in einem semantischen Zusammenhang mit dem Prime (z.B. assoziativ oder kategoriell).
- Syntaktisches Priming: Der geprimte Stimulus und der Prime sind syntaktisch ähnlich (Personen tendieren dazu, die Satztypen zu wiederholen, die bei der Sprachproduktion bereits benutzt wurden; Pickering und Branigan 1999).

- Wortstamm-(Vervollständigungs-)Priming: Der geprimte Stimulus stellt die ersten Buchstaben oder einen Teil des Primes dar (*sag* - *sagte* bzw. *ge-sag-t*).
- Repetition Priming oder Identity Priming: Der geprimte Stimulus und der Prime sind identisch.
- Maskiertes Priming: Der Prime wird nur kurz präsentiert (weniger als 50 Millisekunden), danach folgt eine sogenannte Maske und erst später erscheint der Stimulus, wodurch die Verarbeitung des Primes unterbrochen wird bzw. anders verläuft als ohne Unterbrechung.
- Mediated Priming: Prime und geprimter Stimulus sind durch ein weiteres (vermittelndes) Wort miteinander verbunden, das nicht vorhanden ist bzw. präsentiert wird. Das vermittelnde Wort steht entweder phonologisch oder semantisch mit dem Prime und dem geprimten Stimulus in Zusammenhang.
- Perzeptuelles Priming: Dies betrifft (physische) Formunterschiede zwischen Prime und geprimtem Stimulus und ist modalitätsspezifisch (schriftlich oder auditiv).
- Konzeptuelles Priming: Dies betrifft die Unterschiede der semantischen Verarbeitung der Stimuli, nicht modalitätsspezifisch.
- Cross-modal oder across-modality Priming beschreibt modalitätsübergreifendes Priming.
- Uni-modal oder within-modality Priming beschreibt modalitätsspezifisches Priming.
- Negatives Priming: Manchmal kann die Wiederholung eines Stimulus nicht unbedingt zur Erleichterung der Verarbeitung führen. Wenn die Verarbeitung sich verlangsamt, also gehemmt wird, spricht man von negativem Priming (Tipper 1985, in: Neill 1997).

Im nichtsprachlichen kognitiven Bereich ist es natürlich auch möglich, Priming-Experimente durchzuführen (zum Beispiel Objekt-Priming oder Gesichtserkennungs-Priming), auf diese Bereiche wird hier aber nicht näher eingegangen. Bezüglich der Frage, wie lange der Priming-Effekt andauern kann (mehrere Tage, oder gar ein

Jahr), beachten Sie bitte die weiterführende Literatur bei Tulving et al. (1982) und Cave (1997, beide in: Schacter und Buckner 1998). Priming-Experimente mit Reimen finden Sie beispielsweise bei Blumstein et al. (2000), Bowey (1996), Gordon und Baum (1994), oder Lupker und Williams (1998).

5.2. Das Experiment von Rapp und Samuel (2002)

In dem Experiment von Rapp und Samuel (2002) wurde untersucht, ob der lexikalische Wortabruf bei der Sprachproduktion nur von semantischen oder auch von phonologischen Faktoren abhängt. Dies haben sie in zwei Experimenten überprüft - das erste Experiment war schriftlich gehalten, das zweite mündlich. Beide Experimente wurden auf Englisch durchgeführt.

5.2.1. Aufbau

Die Probanden bekamen Satzpaare mit einer Lücke im zweiten Satz präsentiert und wurden gebeten, eine sinnvolle Ergänzung (ein einzelnes Wort) zu produzieren, das in die Lücke passen würde. Diese mögliche Ergänzung stand semantisch und teilweise auch phonologisch mit den vorhergehenden Sätzen in Zusammenhang. Vor der Lücke befand sich ein Prime. In der Hälfte der Sätze mit dem Prime (im Folgenden als Prime-Item bezeichnet) reimte sich dieser Prime mit einer möglichen Ergänzung durch den Probanden, in der anderen Hälfte nicht. Der sich reimende Prime wurde nicht gesondert markiert, sondern einfach als Teil des Satzes präsentiert. Als Beispiel hierfür wird in der Studie folgender Satz angeführt: „Joe opened the present and hoped he wouldn’t get any clothes. He looked inside the box and found some ____.“

Box ist hier der Prime, der sich auf die mögliche Antwort der Probanden *socks* reimt. In der Kontrollversion dieses Satzes wird der Prime durch ein Wort ersetzt, das in etwa synonym dazu ist: „Joe opened the present and hoped he wouldn’t get any clothes. He looked inside the container and found some ____.“ Der Prime in der Kontrollversion wurde so ausgewählt, dass es sich nicht mit einer möglichen Eingabe des Probanden reimen könnte. Die Autoren stellten die Hypothese auf, dass die Probanden mehr Reimantworten produzieren würden, wenn sie sich mit dem Prime reimen würden, als wenn sie sich nicht mit dem Prime reimen würden.

Experiment 1 war schriftlicher Modalität. Hier wurde weiters der Umstand un-

tersucht, ob die syntaktische Umgebung die Reaktion der Probanden beeinflussen würde. Dazu manipulierten Rapp und Samuel die Sätze so, dass in einer Bedingung der Prime und die Lücke in einem Satz vorkamen (wie die Beispiele oben), in der anderen Bedingung kam der Prime vor der Satzgrenze vor, zum Beispiel: „Joe opened the birthday present and looked inside the box. He found some ____ even though he hoped he wouldn’t get any clothes.“ (Reim-Prime) bzw. „Joe opened the birthday present and looked inside the container. He found some ____ even though he hoped he wouldn’t get any clothes.“ (Nichtreim-Prime). Die Anzahl der Wörter zwischen Prime und Lücke wurde dabei gleich gehalten wie bei der Bedingung innerhalb eines Satzes. Die Autoren stellten die Hypothese auf, dass die Reim-Primes innerhalb der Sätze mehr reimende Vervollständigungen seitens der Probanden hervorrufen sollte, als Nichtreim-Primes.

Das Experiment 2 wurde mündlich gehalten. Dadurch wurde die syntaktische Variable nicht mit hineingenommen, da die Probanden ja nur die Lücke am Ende der Sätze mündlich vervollständigen konnten (und nicht eine Lücke innerhalb der Sätze, bevor sie den gesamten Satz gehört hatten).

5.2.2. Methode

Experiment 1 Rapp und Samuel führten das erste Experiment mit 88 Studenten mit Englisch als L1-Sprache durch. Den Probanden wurden 100 Satzpaare präsentiert, die jeweils kurze Szenarios beschrieben und eine Lücke beinhalteten. Davon waren lediglich 32 Satzpaare „eigentliche“ Items, die restlichen 68 Satzpaare waren Füllitems, damit die Probanden nicht bemerken würden, dass es sich bei der Untersuchung um Reime handelte. Von jedem Satzpaar mit einem Prime gab es vier verschiedene Versionen. Die Variablen waren hier der Prime als Reim oder Nichtreim einer möglichen Eingabe in die Lücke, und die Bedingung, ob der Prime und die Lücke innerhalb eines Satzes oder über die Satzgrenze hinweg vorkommen. Jeweils nur eine Version eines Satzpaars wurde in den Fragebogen inkludiert. Die Füllitems waren immer die gleichen Satzpaare und beinhalteten eine Lücke an einer willkürlich gewählten Stelle des zweiten Satzes. Jeder Proband bekam also 100 Satzpaare, davon waren:

- 8 Satzpaare mit einem Reimwort als Prime innerhalb eines Satzes
- 8 Satzpaare mit einem Reimwort als Prime über die Satzgrenze hinweg

- 8 Satzpaare mit einem Nichtreimwort als Prime innerhalb eines Satzes
- 8 Satzpaare mit einem Nichtreimwort als Prime über die Satzgrenze hinweg
- 68 Füllitems mit einer Lücke an verschiedenen Stellen im zweiten Satz

Somit gab es nur 16 Items bzw. Satzpaare (16%), die tatsächlich einen Reim-Prime beinhalteten. Es wurden 10 Satzpaare pro Seite präsentiert und für jeden Probanden neu randomisiert (sowohl die Seiten als auch die Items innerhalb der Seiten).

Die Probanden wurden gebeten, jedes Satzpaar zu lesen und ein einziges Wort einzugeben, das am besten in den Satz passt. In der Instruktion wurde ein Beispiel gegeben: „We decided to get some fast food for dinner. McDonald’s was only 5 minutes away, and I was in the mood for a ____.“ Beispiele für richtige Antwortmöglichkeiten („hamburger, snack, milkshake“), sowie für falsche Antwortmöglichkeiten (die länger als ein Wort waren, wie „french fries, chicken sandwich“) wurden gegeben.

Experiment 2 Das zweite Experiment wurde mit 61 Studenten durchgeführt. Die Daten einer Person wurden nicht in die Analyse inkludiert, da sie die Instruktionen nicht korrekt befolgt hatte. Experiment 2 stellt die On-line-Version von Experiment 1 dar. Alle Teilnehmer waren Muttersprachler des Englischen. Es wurden vier neue Items mit Reimen hinzugefügt, sodass nun die Itemanzahl der Priming-Sätze 36 betrug. Es gab nur die Bedingung, in der sich die Lücke am Ende des Satzes befand. Die Sätze wurden aufgenommen und den Probanden vorgespielt. 64 Füllitems aus dem ersten Experiment wurden zu den Prime-Items hinzugefügt. Somit hörte jeder Proband 100 Satzpaare, davon waren:

- 18 Satzpaare mit einem Reimwort als Prime
- 18 Satzpaare mit einem Nichtreimwort als Prime
- 64 Füllitems mit einer Lücke am Ende des Satzes

Nur 18% aller Items beinhalteten einen Reim-Prime, der eine Reimantwort hervorrufen könnte. Die Antworten und die Reaktionszeiten der Probanden wurden festgehalten und analysiert.

5.2.3. Ergebnisse

Tatsächlich haben die Probanden mehr Reime nach einem Reim-Prime produziert. In Experiment 1 waren dies 29,8% nach einem Reim-Prime im Vergleich zu 20,1% nach einem Nichtreim-Prime im Mittelwert nach beiden Bedingungen (innerhalb eines Satzes oder über die Satzgrenze hinweg). Diese 20% bedeuten, dass die mögliche Ergänzung der Probanden, die sich mit dem Reim-Prime gereimt hätte, alleine durch den semantischen Kontext auch ohne Reim-Prime hervorgerufen wurde. Es hat sich auch gezeigt, dass mehr Reime produziert wurden, wenn sich der Prime und die Lücke innerhalb eines Satzes befunden haben, als über die Satzgrenze hinweg. Erst im Nachhinein haben die Autoren untersucht, ob die unterschiedliche Orthographie von Prime- und Zielwörtern (*box* und *socks*) einen Einfluss auf die Produktion haben würde, und haben herausgefunden, dass dem nicht so ist.

In Experiment 2 kam ein Reim nach einem Reim-Prime zu 23,4% vor, im Vergleich dazu kam dieses Wort nach einem Nichtreim-Prime nur zu 17,4% vor. Die Probanden produzierten Reimantworten schneller als Antworten, die sich nicht reimten. Dies war signifikant für Teilnehmer, aber nicht für Items.

Aufgrund dieser Ergebnisse sprechen sich Rapp und Samuel für ein interaktives Modell der Wortproduktion aus (siehe Kapitel 4.3.4 *Interaktive Sprachproduktionsmodelle*).

5.3. Adaption des Experiments für das Deutsche

Da das Experiment von Rapp und Samuel (2002) auf Englisch durchgeführt wurde, muss es im Rahmen dieser Arbeit zunächst auf das Deutsche übertragen werden. Dabei sind verschiedene wichtige Dinge zu beachten.

Im Englischen gibt es mehr einsilbige Wörter, die sich reimen, als im Deutschen. Dadurch war es schwierig, passende Items zu finden, die auch einen möglichen semantischen Kontext beschrieben. Im Deutschen gibt es weiters eine andere Satzstellung als im Englischen, sowie trennbare Verben. Um dieses Experiment einheitlich zu gestalten, war es notwendig, die Lücke am Ende des Satzes zu platzieren, was nicht immer die sprachlich einfachste Lösung war.

Im Artikel von Rapp und Samuel (2002) waren im Anhang nur vier Items abgedruckt. Dabei war zu erkennen, dass die Autoren als Prime in diesen vier Items

immer ein Nomen genommen haben. Als potentielle, sich reimende Satzvervollständigung des Probanden war dreimal ein Nomen angeführt, einmal aber auch ein Adverb (*right* im Sinne von *to turn right*). In dem Experiment in meiner Arbeit wurde darauf geachtet, dass die Primes, sowie auch die sich reimende Ergänzung des Probanden nur aus Nomen bzw. Nominalphrasen bestehen würden. Zusätzlich ist zu bemerken, dass bei „to turn ____“ nur zwei richtige Eingaben möglich sind. Bei den Beispielen in dieser Arbeit wurde versucht, bei den Items mit Reimen im Allgemeinen mehr als zwei Möglichkeiten zu geben - nicht nur die potentielle Reimantwort und noch eine Möglichkeit.

Dadurch kommen wir zu einer weiteren Überlegung. Im Englischen war es möglich, den Probanden mitzuteilen, dass sie nur ein Wort eintippen sollten. In einer deutschen Version dieses Experiments ist dies nicht möglich, da es im Deutschen Artikel gibt. Somit muss der Proband auch den Artikel bzw. ein passendes Possessivpronomen eintippen, damit der Satz grammatikalisch korrekt vervollständigt werden kann. Als Autor der Sätze kann man die Artikel nicht vorgeben, da man dadurch die Wortwahl der Probanden bereits vehement einschränken würde. Deswegen sind mehr fehlerhafte bzw. vom Ideal abweichende Antworten seitens der Probanden im Sinne von Tippfehlern, Adjektiveinfügungen, Ausschweifungen, Komposita usw. zu erwarten.

Weiters ist zu beobachten, dass Rapp und Samuel (2002) bei den vier abgedruckten Items bei der Kondition innerhalb der Sätze (wie die Items in meiner Arbeit) die Primes vor eine Konjunktion (*and, when, and, and*) positionierten. Dies wurde in dieser Arbeit ebenso gemacht (*und, doch, sondern*). Im Deutschen konnten daher nur koordinierende Konjunktionen benutzt werden, da aufgrund der Verbendstellung in subordinierten Nebensätzen die Lücke (und somit eine Nominalphrase) nicht am Ende des Satzes hätte stehen können.

5.4. Aufbau des Experiments

Das Experiment hatte Experiment 1 der Studie von Rapp und Samuel (2002) als Grundlage (siehe Kapitel 5.2 *Das Experiment von Rapp und Samuel (2002)*) und ist folgendermaßen aufgebaut: Dem Probanden werden Satzpaare präsentiert, die ein kurzes Szenario beschreiben. Am Ende des zweiten Satzes befindet sich immer eine Lücke, die der Proband füllen soll. Das Experiment wird nur schriftlich durchgeführt.

Von jedem Satzpaar gibt es zwei verschiedene Versionen, eines mit einem Reim-Prime, der sich mit einer potentiellen Satzvervollständigung durch den Probanden reimt, das andere ohne Reim-Prime. Bei dem Satz ohne Reim-Prime wurde darauf geachtet, dass es keine mögliche Satzvervollständigung gibt, die sich auf den Prime reimen würde (siehe das Material in Anhang A).

Im Experiment wird überprüft, ob der Proband eher ein Reimwort als Satzvervollständigung wählt, wenn davor ein Reim-Prime vorkommt, im Vergleich zu der Bedingung ohne Reim-Prime.

Ein großer Unterschied zur Studie von Rapp und Samuel (2002) besteht darin, dass diese nur von einsilbigen Primes und Reimen handelte. In dieser Arbeit wird das Experiment auf zweisilbige Primes und Reime erweitert, um zu überprüfen, ob und inwiefern sich die Ergebnisse auch bei zweisilbigen Reimen zeigen (vergleiche hierzu Kapitel 2.3 *Mehrsilbige Reime*).

In Kapitel 4 *Theoretischer Hintergrund* wurden mehrere Sprachproduktionsmodelle vorgestellt. Wenn L1-Sprecher einer Sprache mehr Reimwörter wählen, wenn davor ein Reim-Prime vorkommt, würde das bedeuten, dass auch die Phonologie einen Einfluss auf die Wortwahl haben würde, nicht nur die Semantik. Dies wäre eine starke Befürwortung von interaktiven Sprachproduktionsmodellen, die ein Rückwirken der Phonologie auf die Semantik postulieren - im Gegensatz zu seriellen Sprachproduktionsmodellen, die zuerst von einer Auswahl auf der Ebene der Semantik ausgehen, bevor die dazugehörige Wortform (Phonologie) aktiviert wird (vgl. 4.3 *Sprachproduktionsmodelle*).

Es sollen die drei in der Einleitung präsentierten Hypothesen überprüft werden:

- Hypothese 1: vergleichbare Ergebnisse wie in der Studie von Rapp und Samuel (2002).
- Hypothese 2: Übertragbarer Effekt auf zweisilbige Reime, wenn auch in eingeschränkterem Umfang.
- Hypothese 3: Mehr unpassende Antworten aufgrund umfangreicherer Antwortmöglichkeiten.

Bei Acheson und MacDonald (2011) (in Kapitel 3.3 *Reime bei sprachgesunden Erwachsenen*) haben wir gesehen, dass es möglich sein kann, dass sich phonologisch

ähnliche Wörter in einem zu engen Rahmen gegenseitig blockieren bzw. hemmen können.

- Hypothese 4 lautet somit, dass die Reimantwort möglicherweise seltener auftreten könnte, wenn davor ein Reim-Prime steht, als bei der Bedingung ohne Reim-Prime (negatives Priming).

5.5. Methode

5.5.1. Probanden

Nicht komplette Datensätze sowie Datensätze von Probanden mit einer anderen Erstsprache als Deutsch wurden nicht mit hineingenommen. Daten von fünf Probanden, die älter waren als 35 Jahre, wurden ebenfalls nicht weiter analysiert, damit eine homogenere Gruppe entstehen würde. Übrig blieben 46 Datensätze von Probanden im Alter zwischen 18 und 35 Jahren (das Durchschnittsalter betrug 25,98 Jahre), davon waren 27 weiblich. Die Probanden kamen aus Österreich und Deutschland.

5.5.2. Material und Design

In der Studie von Rapp und Samuel (2002) wurden den Probanden 100 Satzpaare präsentiert, davon waren jeweils ca. ein Drittel (32% im schriftlichen und 36% im mündlichen Experiment) Prime-Items, sodass die Probanden nicht bemerken würden, dass es sich um Reime handelte. In meinem Experiment im Rahmen dieser Arbeit werden den Probanden nur 60 Satzpaare präsentiert, die Anzahl der Prime-Items beträgt exakt ein Drittel (33%, 20 von 60 Items): 10 Satzpaare mit einsilbigen Primes, 10 Satzpaare mit zweisilbigen Primes, und 40 Füllitems (d.h. Satzpaare), die in Silbenlänge mit den Prime-Items gematcht werden. Die Prime-Items werden von nun an auch als *experimentelle Items* bezeichnet.

Alle Items befinden sich in Anhang A. In Experiment 1 von Rapp und Samuel (2002) befand sich die durch die Probanden zu füllende Lücke nicht nur am Ende des Satzes. Es wurde auch die Variable der Satzgrenze untersucht, wodurch die Lücke bei den Items, sowie bei den Füllitems, an verschiedenen Stellen auftrat. Das Experiment dieser Arbeit verhält sich wie Experiment 2 von Rapp und Samuel (2002): Die Lücken der Satzpaare mit Primes befinden sich nur am Ende des zweiten Satzes - dies ist bei den Füllitems ebenso.

Rapp und Samuel (2002) haben erst im Nachhinein untersucht, ob die unterschiedliche Orthographie der Reimwörter einen Einfluss auf die Produktion hatte (*box - socks*). Sie haben festgestellt, dass dem nicht so war. Dennoch wird in dieser Arbeit darauf geachtet, dass die Reim-Primes und die mögliche Satzvervollständigung die gleiche Orthographie in der Reimkomponente aufweisen würden: Es könnte zum Beispiel das Reimpaar *Tal* und *Schal* genommen werden, nicht aber das Wort *Wahl* in Kombination mit *Tal* oder *Schal*. Weiters wurden für die Sätze mit Reim-Primes nur Wörter verwendet, die eine reguläre Orthographie besitzen.

Bei der Auswahl der Reime wurde darauf geachtet, dass sich keine Komposita in den Reimwörtern oder Reim-Primes befinden, sowie keine morphologische Formen von Wörtern (kein Plural, kein Genitiv) und keine Homophone. Es handelt sich somit um monomorphematische Ein- und Zweisilber. Es sollten keine Redewendungen oder Wortkollokationen auftreten, die die Auswahl eines bestimmten Wortes begünstigen würden (zum Beispiel *eine Frage stellen*, *Schaden erleiden*, jemandes *Hab und Gut* usw.). Auch sollte eine potentielle Reimantwort der Probanden nicht bereits vorher im Satz oder im vorausgegangenen Satz aufgetreten sein (zum Beispiel als Kompositum): Bei der potentiellen Antwort *Reise* sollte beispielsweise davor nicht das Wort *Reisebüro* vorkommen.

Die Reimpaare wurden nach den Wortfrequenzen bzw. Häufigkeitsklassen nach dem Leipziger Wortschatz, der online auf <http://wortschatz.uni-leipzig.de/> verfügbar ist, ausgewählt. Dabei besitzt die mögliche Reimantwort (Lücke) nach einem Reim-Prime immer eine höhere oder gleiche Frequenz (d.h. eine niedrigere Ziffer) als der ihm vorausgegangene Reim-Prime. Ein einzutippendes Reimwort mit niedriger Frequenz ist insofern sehr ungünstig, als dass es umso unwahrscheinlicher wird, dass es als Antwort ausgewählt wird, je niedrigfrequenter es ist (unabhängig von Reim oder Nichtreim-Prime). Vor einem Prime wird kein Adjektiv verwendet. Es tritt nie der gleiche Reim in mehr als einem Item auf. Alle Items wurden im Präteritum verfasst.

Ein Item ist beispielsweise folgendermaßen aufgebaut:

- Sonntag Nachmittag gingen die Großeltern zusammen spazieren. Großvater benutzte seinen Stock (10) und Großmutter trug _____. [einen Rock] (10)
- Sonntag Nachmittag gingen die Großeltern zusammen spazieren. Großvater benutzte *seine Krücke* (16) und Großmutter trug _____.

Die beiden Sätze unterscheiden sich nur in den Nominalphrasen der Primes, die semantisch ähnlich sind. In allen anderen Wörtern sind sie vollkommen ident. Die Zahlen hinter den Primes *Stock*, *Krücke* und der möglichen Reimantwort *Rock* geben ihre Häufigkeitsklasse nach dem Leipziger Wortschatz an. Diese Zahlen, sowie die mögliche Reimantwort der Probanden in eckigen Klammern, befinden sich bei den Items in der Liste im Anhang, werden aber natürlich im Experiment selbst den Probanden nicht gezeigt. Sie dienen lediglich der Itemerstellung.

In den meisten Studien zu Reimen wurde hauptsächlich mit einsilbigen Wörtern gearbeitet, weshalb ich in dieser Arbeit gerne überprüfen würde, ob die vermehrte Produktion von Reimwörtern in einem Reim-Prime-Kontext als in einem Nichtreim-Prime-Kontext nur auf einsilbige Wörter zutrifft, oder auch auf zweisilbige Wörter übertragbar ist. Die beste Variante wäre dabei gewesen, die Variable der ein- und zweisilbigen Primes und Reime so zu testen, dass diese in ein und demselben Satz bzw. Kontext auftreten. Dadurch könnte besser untersucht werden, ob es einen Unterschied zwischen ein- und zweisilbigen Reimwörtern geben würde. Dies war aber aufgrund des inhaltlichen Kontextes bei der Bildung der Items nicht möglich.

Bei den Items mit zweisilbigen Reimen wurden nur solche Wörter verwendet, deren Betonung auf der ersten Silbe liegt, um die phonologische Komponente des *Superreims* anzusprechen (vgl. Abb. 7 auf Seite 17, Berg 1992). Nur ein Superreim kann mehrere Silben umfassen, ein „einfacher“ Reim erstreckt sich immer nur über eine Silbe.

Zusätzlich zu den 40 experimentellen Items gibt es 40 Füllitems, die an die experimentellen Sätze angepasst wurden. Jedes Item der Sätze mit den Reim-Primes hat ein mit ihm in der Silbenanzahl übereinstimmendes Item als Füller. Dies gilt für Version 1 (mit Reim-Prime) und Version 2 (ohne Reim-Primes) eines Items. Es wurden jeweils die Silbenanzahlen in dem Einleitungssatz bis zur Satzgrenze einander angeglichen, sowie die Silbenanzahlen des zweiten Satzes bis zum Prime (inklusive des Primes), und die Silbenanzahlen des restlichen zweiten Satzes bis zur Lücke. Die Lücken befinden sich dabei auch wieder nur am Ende des zweiten Satzes und die Sätze wurden so entworfen, dass Nominalphrasen als Antwort möglich, und sogar wahrscheinlich sind. Konjunktionen der experimentellen Items wurden übernommen: Wo bei den experimentellen Items ein „denn“ vorkam, wurde es auch im Füllsatz benutzt. Somit hat Item 1 der Füllsätze die gleiche Silbenanzahl wie Item 1 der einsilbigen Reime (mit Reim-Prime) usw. Alle Füllitems wurden, wie die experimentellen Items, nur

im Präteritum verfasst. Es kamen hier auch keine Adjektive vor den Prime-Items vor.

5.5.3. Verfahren

Das Experiment wird als Fragebogen mithilfe eines Online-Softwarepaketes entworfen (<https://www.soscisurvey.de/>), sodass die Probanden einfach per Internet darauf zugreifen und den Fragebogen ausfüllen können.

In Anhang A ist der Text zu sehen, der den Probanden die Aufgabe erklären sollte. Zusätzlich sollten sie Angaben zu Geschlecht, Alter, Heimatland, Wohnort, Muttersprache, Bildungsabschluss und Beschäftigung machen. Vor dem Start des Experiments wurden den Probanden drei Übungssitems mit Antwortmöglichkeiten gezeigt, die sich ebenfalls im Anhang befinden.

Das eigentliche Experiment geht folgendermaßen von statten: Den Probanden werden die 40 Füllitems präsentiert, sowie jeweils Version 1 oder Version 2 eines experimentellen Items (Reim-Prime oder Nichtreim-Prime). Somit sieht jeder Proband 5 Items mit Reim-Prime und 5 Items ohne Reim-Prime pro einsilbiger oder zweisilbiger Reimbedingung. Das Erscheinen von Version 1 oder Version 2 geschieht randomisiert mit dem sogenannten „Urnen-Verfahren“ (siehe https://www.soscisurvey.de/help/doku.php/de:create:random_urns), das eine gleichmäßige Verteilung gewährleisten soll. So kann es nicht geschehen, dass eine Version viel häufiger angeboten wird als die andere. Die Füllitems sind in jedem Fragebogen die gleichen. Da die Gesamtanzahl der eines Probanden präsentierten Items 60 beträgt (40 Füllitems und 20 experimentelle Items), werden über 6 Seiten hinweg jeweils 10 Items präsentiert. Die Items werden über alle 6 Seiten hinweg randomisiert. Somit liest jeder Proband:

- 5 Satzpaare mit einem einsilbigen Reimwort als Prime
- 5 Satzpaare mit einem einsilbigen Nichtreimwort als Prime
- 5 Satzpaare mit einem zweisilbigen Reimwort als Prime
- 5 Satzpaare mit einem zweisilbigen Nichtreimwort als Prime
- 40 Füllitems mit einer Lücke am Ende des Satzes

5.6. Ergebnisse

Bei fünf Satzpaaren kam das gesuchte Wort (das sich auf den Prime in der ersten Version reimen sollte) weder in Version 1 (mit dem Reim-Prime), noch in Version 2 (mit dem Nichtreim-Prime) vor. Es handelt sich hierbei (siehe Anhang A) bei den einsilbigen Reimen um Item 2 (Reim-Prime: *Wicht*, Reim-Zielitem: *Licht*), 5 (*Heer* - *Speer*), 8 (*Schach* - *Fach*), 9 (*Fee* - *Klee*), bei den zweisilbigen Reimen um Item 5 (*Puppe* - *Suppe*). Es könnte verschiedene Gründe dafür geben, warum keiner der Probanden auf das gesuchte Wort kam, was im folgenden Kapitel (Kapitel 6 *Diskussion und Zusammenfassung*) näher beleuchtet werden soll. Die 5 oben genannten Items wurden aus der weiteren Analyse herausgenommen - somit blieben 15 weiter analysierbare Items übrig.

Bei der Version mit dem Reim-Prime kam das sich reimende Wort zu 26,94% vor, bei der Version mit dem Nichtreim-Prime kam das gleiche Wort (das sich mit dem Reim-Prime reimen würde) zu 25,14% vor. Es zeigte sich somit kein signifikanter Unterschied in den beiden Versionen bzw. kein Reim-Priming. Wenn man die einsilbigen und die zweisilbigen Reime einzeln betrachtet, zeigt sich bei den einsilbigen Reimen eine generell höhere Anzahl der verwendeten Reimwörter (29,64% mit Reim-Prime, 28,96% mit Nichtreim-Prime) im Vergleich zu den zweisilbigen Reimen (25,14% mit Reim-Prime, 22,59% mit Nichtreim-Prime). Man kann somit sagen, dass das gesuchte Wort bei den einsilbigen Reimen generell öfter als Antwort kam (dies trifft auf beide Versionen eines Items zu, nur 0,68% Unterschied), während bei den zweisilbigen Reimen ein etwas größerer Unterschied zwischen Version 1 und Version 2 eines Items bestand (2,55% Unterschied).

6. Diskussion und Zusammenfassung

Die im Experiment gewonnenen Ergebnisse sprechen gegen einen Einfluss von Reimen in der Sprachproduktion bzw. beim Abruf von Wörtern. Die Hypothesen 1, 2 und 4 wurden somit nicht bewiesen. Hypothese 3 wurde bewiesen, da die Probanden tatsächlich häufiger Adjektive einfügten, Komposita bildeten oder gar halbe Sätze eintippten. Die Resultate könnten verschiedene Ursachen haben, und zum Teil sprachspezifischer Natur sein: Eine Möglichkeit wäre der Umstand, dass in diesem Experiment - da es auf Deutsch durchgeführt wurde - ein Nomen und der dazugehörige Artikel eingefügt werden mussten (anstatt im Englischen nur das Nomen). Dadurch kam es häufig zu Formen wie beispielsweise „zu der“, statt zur üblicheren Version „zur“, da *zu* bereits vorgegeben war, *der* noch einge tippt werden musste. Dies war für die Probanden eventuell ungewöhnlich, wie sie auch am Ende des Fragebogens manchmal angemerkt haben und könnte die Eingabe beeinflusst haben. Der Fokus konnte nicht auf das Nomen alleine gelegt werden, die Einfügung des Artikels könnte als Störung empfunden worden sein und die Probanden daran gehindert haben, spontaner ihre Ideen einzutippen.

Ein weiterer Faktor waren Komposita: Im Englischen gibt es weniger Komposita als im Deutschen, und die deutschen Probanden waren versucht, des Öfteren Komposita einzufügen. Am Anfang des Fragebogens wurden sie aber darum gebeten, nur einfache Wörter zu benutzen, was das Ausfüllen des Fragebogens verkomplizierte, siehe Anhang A.5: *Bitte benutzen Sie, wenn möglich, **keine** zusammengesetzten Wörter, wie: die Eintrittskarte (Eintritts-karte)*. Ich habe die Probanden darum gebeten, keine Komposita zu verwenden, da die Analyse ja nur ein- und zweisilbige Reime betreffen sollte und auch bei den Primes keine Komposita vorkamen. Aufgrund der Zusammensetzung zweier Wörter besitzt in Komposita jedes der Wörter einzeln gesehen einen anderen Rhythmus, als wenn es alleine stehen würde. Es wäre

aber interessant, eine ähnliche Untersuchung ohne Hinweis auf Nichtverwendung von Komposita durchzuführen und die Ergebnisse zu vergleichen.

Weiters war vielleicht die Anzahl der Items (10 Reim-Primes, die eine Person tatsächlich gesehen hat, davon 5 zweisilbig) zu gering, um einen Effekt unter den oben genannten Umständen (Einfügen von zwei Wörtern etc.) zu erkennen. Möglicherweise spielte auch die Frequenz der Wörter eine gewisse Rolle und war zu gering, um einen Reimeffekt erkennen zu lassen.

Es wäre auch denkbar, dass die Auswahl und der Aufbau der Items nicht einwandfrei gelungen ist oder das semantische Szenario nicht dazu führte, dass das Zielitem (phonologisch oder gar semantisch) abgerufen wurde. Ein Indiz dafür wäre die Tatsache, dass bei fünf Items das gesuchte Reimwort weder in der Bedingung mit dem Reim-Prime, noch in der Bedingung mit dem Nichtreim-Prime vorgekommen ist und somit als zu abwegig oder gar als zu selten bzw. als nicht frequent genug bezeichnet werden kann. Interessant ist hierbei, dass bei den zweisilbigen Reimen bei nur einem Item das gesuchte Reimwort in beiden Bedingungen nie vorgekommen ist.

Obwohl bei der Erstellung der Items darauf geachtet wurde, dass die Reimkomponente der Reim-Primes und der Reim-Zielitems auch im Schriftbild identisch war, wäre es möglich, dass das Experiment andere Ergebnisse erzielen würde, wenn es auditiv durchgeführt werden würde. Es wäre einen Versuch wert, das Schriftbild der Reim-Primes und Reim-Zielitems außer Acht zu lassen, und stattdessen Wörter zu benutzen, die eine höhere Frequenz besitzen.

Man kann zusammenfassend sagen, dass es möglicherweise sprachspezifische Unterschiede gibt, die teilweise zu einer anderen Itemauswahl und zu anderen Ergebnissen geführt haben, die sich mit den Ergebnissen von Rapp und Samuel (2002) nicht decken.

Rapp und Samuel (2002) befürworteten ein interaktives Sprachproduktionsmodell und meinten, dass es realitätsgetreuer wirkte als ein serielles Sprachproduktionsmodell, da die Phonologie die Produktion der Wörter beeinflusste. In meinem Experiment war dieser Effekt nicht zu erkennen. Dennoch bin ich der Meinung, dass ein interaktives Sprachproduktionsmodell, und damit die Einwirkung von späteren Verarbeitungsschritten auf frühere, mehr Sprachverarbeitungsprozesse beschreiben kann als ein serielles Modell.

Mit dieser Arbeit habe ich versucht, einen Einblick in Reime von verschiedenen Blickwinkeln aus zu geben. Nach einer Einführung in ihren phonologischen Aufbau

wurde auf die Grenzen von Reimen, wie sie in der wissenschaftlichen Literatur behandelt werden, eingegangen. Oft wird von einem Reim nur als einsilbiges Wort gesprochen. Wie wir gesehen haben, kann er sich aber auch über mehrere Silben erstrecken. Vielleicht wäre es sinnvoll, weitere Untersuchungen mit zwei- oder mehrsilbigen Reimen durchzuführen, um die Bandbreite von Reimen zu erforschen, sowie die Frage, ob mehrsilbige Reime anders verarbeitet werden als einsilbige. Wenn es mir möglich war, habe ich in den einzelnen Unterkapiteln des aktuellen Forschungsstands auf Studien mit mehrsilbigen Reimen verwiesen, die meistens vor noch nicht allzu langer Zeit veröffentlicht wurden. Dies könnte darauf hinweisen, dass die Forschung in diese Richtung geht.

Kapitel 4 sollte die Grundlage und den theoretischen Hintergrund für das Experiment liefern, und es wurden darin Themen wie das mentale Lexikon, die Wortwahl, und der Einfluss von Orthographie und Phonologie behandelt. Es ist wahrscheinlich, dass das graphematische und das phonologische Lexikon zwei unterschiedliche Lexika darstellen, woraus folgt, dass die Schriftsprache nicht zwingend von der Phonologie abhängen muss, aber kann (und umgekehrt). Wie gesagt, würde das Experiment vielleicht andere Ergebnisse erzielen, wenn es auditiv durchgeführt werden würde. Auch in der Studie von Rapp und Samuel (2002) wurde bei einer auditiven Version ihres Experiments ein signifikanter Effekt für Reim-Priming gefunden, wenn auch nicht so viele Reimwörter vorkamen wie in der schriftlichen Version.

In jedem Fall scheint das Thema Reime in Zusammenhang mit Mehrsilbigkeit, Einfluss auf die Wortwahl, und womöglich sogar Rhythmus ein vielversprechendes Thema für die Zukunft zu sein und erfordert weitere Forschung.

7. Bibliographie

ACHESON, Daniel J./MACDONALD, Maryellen C. (2011), „The rhymes that the reader perused confused the meaning: Phonological effects during on-line sentence comprehension“, *Journal of Memory and Language*, 65, 193-207.

AITCHISON, Jean/STRAF, Miron (1981), „Lexical storage and retrieval: a developing skill?“, *Linguistics*, 19, 751-795.

ALARIO, François-Xavier/SCHILLER, Niels O./DOMOTO-REILLY, Kimiko/CARMAZZA, Alfonso (2003), „The role of phonological and orthographic information in lexical selection“, *Brain and Language*, 84, 372-398.

ALTMANN, Hans/ZIEGENHAIN, Ute (2002), *Phonetik, Phonologie und Graphemik fürs Examen*, Wiesbaden: Westdeutscher Verlag GmbH, 168pp.

BERG, Thomas (1989), „On the internal structure of polysyllabic monomorphemic words: The case for superrimes“, *Studie Linguistica*, 43, 5-32.

BERG, Thomas (1992), „Umriss einer psycholinguistischen Theorie der Silbe“, in: EISENBERG, Peter/RAMERS, Karl Heinz/VATER, Heinz (eds.), *Silbenphonologie des Deutschen. Studien zur deutschen Grammatik*, 42, Tübingen: Gunter Narr Verlag, 45-99.

BLUMSTEIN, Sheila E./MILBERG, William/BROWN, Todd/HUTCHINSON, Adele/KUROWSKI, Kathleen/BURTON, Martha W. (2000), „The Mapping from Sound Structure to the Lexicon in Aphasia: Evidence from Rhyme and Repetition Priming“, *Brain and Language*, 72, 75-99.

BOCK, Kathryn/LEVELT, Willem J. L. (1994), „Language Production: Grammatical encoding“, in: GERNSBACHER, M. (ed.), *Handbook of psycholinguistics*, San Diego, CA: Academic Press, 945-984.

BONIN, Patrick/FAYOL, Michel/GOMBERT, Jean-Emile (1998), „An Experimental Study of Lexical Access in the Writing and Naming of Isolated Words“, *International Journal of Psychology*, 33/4, 269-286.

BOWEY, Judith A. (1996), „Phonological Recoding of Nonword Orthographic Rhyme Primes“, *Journal of Experimental Psychology: Learning, Memory and Cognition*, 22/1, 117-131.

BRYANT, Peter E./BRADLEY, Lynette/MACLEAN, Morag/CROSSLAND, John (1989), „Nursery rhymes, phonological skills and reading“, *Journal of Child Language*, 16, 407-428.

BURTON, Martha Whipple (1989), *Associative, mediated, and rhyme priming: A study of lexical processing*, Unpublished doctoral dissertation, Brown University.

BUßMANN, Hadumod (ed.) (⁴2008), *Lexikon der Sprachwissenschaft*, Stuttgart: Alfred Kröner Verlag, 816pp.

CARAMAZZA, Alfonso (1997), „How Many Levels of Processing Are There in Lexical Access?“, *Cognitive Neuropsychology*, 14/1, 177-208.

CHOLEWA, Jürgen/MANTEY, Stefanie (2007), *Grammatische Grundlagen für die Sprachtherapie*, München: Elsevier GmbH, 165pp.

COCH, Donna/GROSSI, Giordana/SKENDZEL, Wendy/NEVILLE, Helen (2005), „ERP Nonword Rhyming Effects in Children and Adults“, *Journal of Cognitive Neuroscience*, 17/1, 168-182.

COCH, Donna/HART, Tory/MITRA, Priya (2008), „Three kinds of rhymes: An ERP study“, *Brain and Language*, 104, 230-243.

DAVIS, Stuart (1989), „On a non-argument for the rhyme“, *Journal of Linguistics*, 25, 211-217.

DE BLESER, Ria/CHOLEWA, Jürgen/STADIE, Nicole/TABATABAIE, Sia (2004), *LEMO Lexikon modellorientiert. Einzelfalldiagnostik bei Aphasie, Dyslexie und Dysgraphie*, München: Elsevier GmbH, 137pp.

DELL, Gary S. (1986), „A Spreading-Activation Theory of Retrieval in Sentence Production“, *Psychological Review*, 93/3, 283-321.

DELL, Gary S. (1988), „The Retrieval of Phonological Forms in Production: Tests of Predictions from a Connectionist Model“, *Journal of Memory and Language*, 27, 124-142.

DELL, Gary S./O'SEAGHDHA, Padraig G. (1991), „Mediated and Convergent Lexical Priming in Language Production“, *Psychological Review*, 98/4, 604-614.

DELL, Gary S./SCHWARTZ, Myrna F./MARTIN, Nadine/SAFFRAN, Eleanor M./GAGNON, Deborah A. (1997), „Lexical Access in Aphasic and Nonaphasic Speakers“, *Psychological Review*, 104/4, 801-838.

DUDEENREDAKTION (⁸2009), *Die Grammatik. Unentbehrlich für richtiges Deutsch*, Mannheim u.a.: Dudenverlag, §§1-196 (= Duden 4) [Autor des Kapitels: Peter EISENBERG], 1343pp.

DUNCAN, Lynne G./SEYMOUR, Philip H. K. (2000), „Phonemes and rhyme in the development of reading and metaphonology: The Dundee longitudinal study“, in: BADIAN, Nathalie A. (ed.), *Prediction and Prevention of reading failure*, Parkton, MD: The York Press, 199-221.

DUNCAN, Lynne G./SEYMOUR, Philip H. K./BOLIK, Fiona (2007), „Rimes and superrimes: An exploration of children's disyllabic rhyming skills“, *British Journal of Psychology*, 98, 199-221.

EYSENCK, Michael W./KEANE, Mark T. (⁶2010), *Cognitive Psychology. A Student's Handbook*, Hove/New York: Psychology Press.

FRAWLEY, William J. (ed.) (²2003), *International Encyclopedia of Linguistics*, New York: Oxford University Press Inc.

GLASSPOOL, David W./HOUGHTON, George (2005), „Serial order and consonant-vowel structure in a graphemic output buffer model“, *Brain and Language*, 94, 304-330.

GLÜCK, Helmut (ed.) (⁴2010), *Metzler Lexikon Sprache*, Stuttgart: J.B. Metzler'sche Verlagsbuchhandlung und Carl Ernst Poeschel Verlag GmbH, 814pp.

GORDON, Jean K./BAUM, Shari R. (1994), „Rhyme Priming in Aphasia: The Role of Phonology in Lexical Access“, *Brain and Language*, 47, 661-683.

GOSWAMI, Usha/MEAD, Felicity (1992), „Onset and rime awareness and analogies in reading“, *Reading Research Quarterly*, 27/2, 152-162.

GROSSI, Giordana/COCH, Donna/COFFEY-CORINA, Sharon/HOLCOMB, Phillip J./NEVILLE, Helen J. (2001), „Phonological Processing in Visual Rhyming: A Developmental ERP Study“, *Journal of Cognitive Neuroscience*, 13/5, 610-625.

HALL, Tracy Alan (1992), *Syllable Structure and Syllable-Related Processes in German*, Tübingen: Max Niemeyer Verlag, 246pp.

HALL, Tracy Alan (2000), *Phonologie. Eine Einführung*, Berlin/New York: Walter de Gruyter GmbH, 360pp.

HARLEY, Trevor A. (1984), „A Critique of Top-down Independent Levels Models of Speech Production: Evidence from Non-plan-Internal Speech Errors“, *Cognitive Science*, 8, 191-219.

HENSON, Richard N.A. (2003), „Neuroimaging studies of priming“, *Progress in Neu-*

robiology, 70, 53-81.

HOUGHTON, George/ZORZI, Marco (2003), „Normal and impaired spelling in a connectionist dual-route architecture“, *Cognitive Neuropsychology*, 20/2, 115-162.

HUBER, Walter/POECK, Klaus/SPRINGER, Luise (2006), *Klinik und Rehabilitation der Aphasie. Eine Einführung für Therapeuten, Angehörige und Betroffene*, Stuttgart: Georg Thieme Verlag, 162pp.

HYESUK, Cho (2011), „Exploring a common neural substrate of reading and spelling“, *Dissertation Abstracts International, Section A: The Humanities and Social Sciences*, 71/11, University of Arizona, 4001pp.

INKELAS, Sharon (2003), „J's rhymes: a longitudinal case study of language play“, *Journal of Child Language*, 30, 557-581.

JESCHENIAK, Jorg D./LEVELT, Willem J. L. (1994), „Word frequency effects in speech production: Retrieval of syntactic information and of phonological form“, *Journal of Experimental Psychology: Learning, Memory and Cognition*, 20, 824-843.

JESPERSEN, Otto (1904), *Phonetische Grundfragen*, Harvard University: B.G. Teubner, 185pp.

KAY, Janice/LESSER, Ruth/COLTHERT, Max (1992), *PALPA: Psycholinguistic Assessments of Language Processing in Aphasia*, Hove, UK: Lawrence Erlbaum Associates Ltd.

KIRTLEY, Clare/BRYANT, Peter/MACLEAN, Morag/BRADLEY, Lynette (1989), „Rhyme, Rime and the Onset of Reading“, *Journal of Experimental Child Psychology*, 48, 224-245.

KLICPERA, Christian/GASTEIGER-KLICPERA, Barbara (1994), „Die langfristige Entwicklung der mündlichen Lesefähigkeit bei schwachen und guten Lesern“, *Zeitschrift für Entwicklungspsychologie und Pädagogische Psychologie*, 26 278-290.

KOHLER, Klaus J. (²1995), *Einführung in die Phonetik des Deutschen* (Grundlagen der Germanistik, 20), Berlin: Erich Schmidt Verlag GmbH, 249pp.

KREINER, David S. (1992), „Reaction Time Measures of Spelling: Testing a Two-Strategy Model of Skilled Spelling“, *Journal of Experimental Psychology: Learning, Memory and Cognition*, 18/4, 765-776.

LEVELT, Willem J.L./Roelofs, Ardi/Meyer, Antje S. (1999), „A theory of lexical access in speech production“, *Behavioral and Brain Sciences*, 22, 1-75.

LIBBEN, Gary/Jarema, Gonia (2002), „Mental Lexicon Research in the New Millennium“, *Brain and Language*, 81, 2-11.

LUPKER, Stephen J./WILLIAMS, Bonnie A. (1989), „Rhyme Priming of Pictures and Words: A Lexical Activation Account“, *Journal of Experimental Psychology: Learning, Memory and Cognition*, 15/6, 1033-1046.

MAAS, Utz (2006), *Phonologie. Einführung in die funktionale Phonetik des Deutschen*, Göttingen: Vandenhoeck & Ruprecht GmbH & Co. KG, 392pp.

MCGLONE, Matthew S./TOFIGHBAKHS, Jessica (2000), „BIRDS OF A FEATHER FLOCK CONJOINTLY (?): Rhyme as Reason in Aphorisms“, *Psychological Science*, 11/5, 424-428.

MEYER, Antje S. (1990), „The time course of phonological encoding in language production: The encoding of successive syllables of a word“, *Journal of Memory and Language*, 29, 524-545.

MICELI, Gabriele/BENVEGNÙ, Barbara/CAPASSO, Rita/CARAMAZZA, Alfonso (1997), „The Independence of Phonological and Orthographic Lexical Forms: Evidence from Aphasia“, *Cognitive Neuropsychology*, 14/1, 35-69.

NEILL, W. Trammell (1997), „Episodic Retrieval in Negative Priming and Repetition Priming“, *Journal of Experimental Psychology: Learning, Memory, and Cognition*,

23/6, 1291-1305.

NICKELS, Lindsay/BEST, W (1996), „Therapy for naming deficits (part I): Principles, puzzles and progress“, *Aphasiology*, 10, 21-47.

NICKELS, Lindsay/COLE-VIRTUE, Jennifer (2004), „Reading tasks from PALPA: How do controls perform on visual lexical decision, homophony, rhyme and synonym judgements?“, *Aphasiology*, 18/2, 103-126.

NOACK, Christina (2010), *Phonologie*, Heidelberg: Universitätsverlag Winter GmbH, 100pp.

PERRIN, Fabien/GARCIA-LARREA, Luis (2003), „Modulation of the N400 potential during auditory phonological/semantic interaction“, *Cognitive Brain Research*, 17, 36-47.

PICKERING, Martin J./BRANIGAN, Holly P. (1999), „Syntactic Priming in language production“, *Trends in Cognitive Sciences*, 3/4, 136-141.

PTOK, Martin/BERENDES, Karin/GOTTAL, Stephanie/GRABHERR, Britta/SCHNEEBERG, Jennifer/WITTLER, Marion (2008), „Phonologische Verarbeitung“, *Monatsschrift Kinderheilkunde*, 156/9, 860-866.

RAPP, Brenda/BENZING, Lisa/CARAMAZZA, Alfonso (1997), „The Autonomy of Lexical Orthography“, *Cognitive Neuropsychology*, 14/1, 71-104.

RAPP, Brenda/GOLDRICK, Matthew (2000), „Discreteness and Interactivity in Spoken Word Production“, *Psychological Review*, 107/3, 460-499.

RAPP, David N./SAMUEL, Arthur G. (2002), „A Reason to Rhyme: Phonological and Semantic Influences on Lexical Access“, *Journal of Experimental Psychology: Learning, Memory and Cognition*, 28/3, 564-571.

RAYMER, Anastasia M./THOMPSON, Cynthia K./JACOBS, Beverly J./LE GRAND,

Henriette R. (1993), „Phonological treatment of naming deficits in aphasia: model-based generalization analysis“, *Aphasiology*, 7/1, 27-53.

RING, Jeremiah James (2001), „An Investigation of Orthographic Representation in Skilled Spelling Production and Word Recognition“, *Dissertation Abstracts International, Section B: Sciences and Engineering*, 61/12, University of Colorado, Boulder, 6731pp.

ROELOFS, Ardi (1992), „A spreading-activation theory of lemma retrieval in speaking“, *Cognition*, 42, 107-142.

RUGG, Michael D. (1984a), „Event-related potentials and the phonological processing of words and non-words“, *Neuropsychologia*, 22, 435-443.

RUGG, Michael D. (1984b), „Event-Related Potentials in Phonological Matching Tasks“, *Brain and Language*, 23, 225-240.

RUGG, Michael D./BARRETT, Sarah E. (1987), „Event-Related Potentials and the Interaction between Orthographic and Phonological Information in a Rhyme-Judgement Task“, *Brain and Language*, 32, 336-361.

SCHACTER, Daniel L./Buckner, Randy L. (1998), „Priming and the Brain“, *Neuron*, 20, 185-195.

SCHWARZ, Monika (³2008),, *Einführung in die Kognitive Linguistik*, Tübingen: Francke, 298pp.

SEYMOR, Philip H. K./Evans, Henryka M. (1994), „Levels of phonological awareness and learning to read“, *Reading and Writing*, 6, 221-250.

SHELTON, Jennifer R./Weinrich, Michael (1997), „Further Evidence of a Dissociation between Output Phonological and Orthographic Lexicons: A Case Study“, *Cognitive Neuropsychology*, 14/1, 105-129.

SPENCER, Kristie A./DOYLE, Patrick J./MCNEIL, Malcolm R./WAMBAUGH, Julie L./PARK, Grace/CARROLL, Brian (2000), „Examining the facilitative effects of rhyme in a patient with output lexicon damage“, *Aphasiology*, 15/5,6, 567-584.

STADIE, Nicole/CHOLEWA, Jürgen/DE BLESER, Ria/TABATABAIE, Sia (1994), „Das neurolinguistische Expertensystem LeMo. I. Theoretischer Rahmen und Konstruktionsmerkmale des Testteils LEXIKON“, *Neurolinguistik*, 8/1, 1-25.

STADIE, Nicole/KEIM, Roland/DE BLESER, Ria (2003), „Aufgaben zur Überprüfung phonologischen Wissens/phonologischer Bewußtheit (PhoWi). Theoretischer Rahmen und Konstruktionsmerkmale“, *Neurolinguistik*, 17/2, 135-160.

STAFFELDT, Sven (2010), *Einführung in die Phonetik, Phonologie und Graphematik des Deutschen. Ein Leitfaden für den akademischen Unterricht*, Tübingen: Stauffenburg Verlag Brigitte Narr GmbH, 191pp.

TREIMAN, Rebecca (1992), „The role of intrasyllabic units in learning to read and spell“, in: GOUGH, Philip B./EHRI, Linnea C./TREIMAN, Rebecca (eds.), *Reading Acquisition*, Hillsdale, NJ: Erlbaum, 85-106.

TREIMAN, Rebecca/ZUKOWSKI, Andrea (1996), „Children’s sensitivity to syllables, onsets, rimes, and phonemes“, *Journal of Experimental Child Psychology*, 61, 193-215.

TWADDELL, William F. (1939), „Combinations of consonants in stressed syllables in German“, *Acta Linguistica*, 1, 189-199.

TWADDELL, William F. (1940), „Combinations of consonants in stressed syllables in German (continued)“, *Acta Linguistica*, 2, 31-50.

VATER, Heinz (1992), „Zum Silben-Nukleus im Deutschen“, in: EISENBERG, Peter/RAMERS, Karl Heinz/VATER, Heinz (eds.), *Silbenphonologie des Deutschen. Studien zur deutschen Grammatik*, 42, Tübingen: Gunter Narr Verlag, 100-133.

7. Bibliographie

VATER, Heinz (⁴2002), *Einführung in die Sprachwissenschaft*, München: Wilhelm Filk Verlag, 336pp.

WAGENSVELD, Barbara/SEGERS, Eliane/VAN ALPHEN, Petra/HAGOORT, Peter/VERHOEVEN, Ludo (2012), „A neurocognitive perspective on rhyme awareness: The N450 rhyme effect“, *Brain Research*, 1483, 63-70.

WAGENSVELD, Barbara/SEGERS, Eliane/VAN ALPHEN, Petra/VERHOEVEN, Ludo (2013), „The role of lexical representations and phonological overlap in rhyme judgements of beginning, intermediate and advanced readers“, *Learning and Individual Differences*, 23, 64-71.

WEBER-FOX, Christine/SPENCER, Rebecca/CUADRADO, Elizabeth/SMITH, Anne (2003), „Development of Neural Processes Mediating Rhyme Judgements: Phonological and Orthographic Interactions“, *Developmental Psychobiology*, 43, 128-145.

WELDON, Mary Susan (2000), „Mechanisms Underlying Priming on Perceptual Tests“, *Journal of Experimental Psychology Learning, Memory and Cognition*, 17/3, 526-541.

WIESE, Richard (²2006), *The phonology of German*, New York: Oxford University Press Inc., 358pp.

WISENBURN, Bruce/MAHONEY, Kate (2009), „A meta-analysis of word-finding treatments for aphasia“, *Aphasiology*, 23/11, 1338-1352.

Internetquelle

Leipziger Wortschatz, URL: <http://wortschatz.uni-leipzig.de/>, abgerufen in Dezember 2012 und Jänner 2013

8. Abbildungsverzeichnis

1.	Möglichkeiten der Verteilung von zwei Konsonanten im Onset (Duden 2009, 8. Auflage, S. 42)	7
2.	Silbenstruktur (Staffeldt 2010, S. 134)	10
3.	Die Sonoritätshierarchie (Duden 2009, 8. Auflage, S. 40)	13
4.	Die Sonoritätshierarchie nach Jespersen (1904, in: Noack 2010, S. 60)	14
5.	Die Sonoritätshierarchie nach Maas (2006, S. 121)	14
6.	Silbenstruktur nach Treiman (1992, in: Duncan et al. 2007, S. 201)	17
7.	Die hierarchische Organisation von Silben in mehrsilbigen Wörtern (Berg 1992, S. 80)	17
8.	Silbenstruktur nach Berg (1989, in: Duncan et al. 2007, S. 202)	18
9.	Reaktionszeit und Fehlerrate beim Reimentscheiden bei gesunden Erwachsenen (Nickels und Cole-Virtue 2004, S. 113)	26
10.	Reaktionszeit beim lexikalischen Entscheiden bei gesunden Erwachsenen (Blumstein et al. 2000, S. 83)	27
11.	Reaktionszeit beim lexikalischen Entscheiden bei Broca- und Wernicke-Asphasikern (Blumstein et al. 2000, S. 84)	32
12.	Anzahl der korrekten mündlichen Benennungen der trainierten und untrainierten Items in Prozent (Spencer et al. 2000, S. 575)	37
13.	(A) Übersicht der EKPs bei allen drei Bedingungen. (B) Topographische Karten der drei Bedingungen und Differenzdarstellung (Wagensveld et al. 2012, S. 66)	42
14.	Das Logogenmodell in LeMo in Anlehnung an Patterson 1988 (De Bleser et al. 2004, S. 7)	48
15.	Das <i>Independent Network model</i> von Caramazza (1997, S. 196)	55

16.	Das <i>Independent Network model</i> (detaillierter) von Caramazza (1997, S. 197)	56
17.	Das <i>Standard Top-Down Serial Processing Model of Speech Production</i> (Harley 1984, S. 193)	57
18.	Sprachproduktionsmodell (seriell) von Levelt et al. 1999, S. 3	59
19.	Fragment des lexikalischen Netzwerks beim Zugriff aufs Lexikon (Levelt et al. 1999, S. 4)	62
20.	Lexikalisches Netzwerk bei der Sprachproduktion (Dell 1986, S. 290) .	65
21.	Das lexikalische und das Wortformnetzwerk (Dell 1988, S. 128)	68
22.	Lexikalisches Netzwerk bei der Sprachproduktion (Dell und O'Seaghdha 1997, S. 805)	69
23.	Darstellung der phonologischen und orthographischen Repräsentationen (Caramazza 1997, S. 190)	70
24.	Das <i>dual-route model for spelling</i> von Houghton und Zorzi (2003, S. 117)	73

9. Anhang

A. Material (Experiment)

In der a)-Version befindet sich der Satz mit dem Reim-Prime, die b)-Version ist der Satz ohne Reim-Prime mit einem semantisch ähnlichen Wort zu dem Reim-Prime *in kursiver Schrift*. Die potentielle Reimantwort der Probanden befindet sich in eckiger Klammer hinter dem Satz. Hinter den Prime-Wörtern und der potentiellen Reimantwort befindet sich in Klammer die jeweilige Häufigkeitsklasse dieser Wörter (Leipziger Wortschatz).

Etwaige Überlegungen zur Auswahl des Kontrollitems habe ich nach dem Item eingefügt. Oft habe ich das Nichtprime-Wort so gewählt, dass kein Reim als Eingabe in die Lücke möglich ist, unabhängig davon, ob es semantisch nicht sowieso ganz unpassend wäre.

A.1. Items mit einsilbigen Reimen

Die Itemanzahl für Satzpaare mit einsilbigen Reimen soll 10 betragen ($n = 10$).

1. a) Sofie ging am Samstag in ein Gruselkabinett. Es war stockdunkel, sie berührte die Wand (10) und fühlte plötzlich _____. [eine Hand] (8)
b) Sofie ging am Samstag in ein Gruselkabinett. Es war stockdunkel, sie berührte *den Vorhang* (12) und fühlte plötzlich _____.

Mauer konnte ich nicht nehmen, da sich *Hauer*, *Schauer* und *Bauer* darauf reimen. *Widerstand* war auch nicht passend, da es sich auf *-and* reimt.

2. a) Im Feenland traf Peter viele interessante Gestalten. Zum Beispiel sah er einmal einen Wicht (16) und dieser trug _____. [ein Licht] (9)
b) Im Feenland traf Peter viele interessante Gestalten. Zum Beispiel sah er einmal *einen Gnom* (18) und dieser trug _____.

Einen Zwerg habe ich nicht genommen, da es sich auf *Berg* reimt, und in einem Feenland ist es vielleicht möglich, einen Berg zu tragen. Ein *Strom* ist schon um einiges schwieriger zu tragen (Reim auf *Gnom*).

3. a) Karl freute sich sehr über die Gehaltserhöhung. Er erhielt den Lohn (11) und kaufte sofort ein Spielzeug für _____. [seinen/den Sohn] (8)
- b) Karl freute sich sehr über die Gehaltserhöhung. Er erhielt *das Honorar* (13) und kaufte sofort ein Spielzeug für _____.

Das Gehalt und *Das Geld* habe ich nicht genommen, da sich diese Wörter mit zu vielen anderen Wörtern reimen (*Wald, Held, Feld* etc.).

4. a) Der königliche Ball war ein gesellschaftliches Großereignis. Das Schloss erstrahlte in seinem Glanz (11) und der Thronfolger eröffnete _____. [den Tanz] (11)
- b) Der königliche Ball war ein gesellschaftliches Großereignis. Das Schloss erstrahlte *in seiner Schönheit* (13) und der Thronfolger eröffnete _____.

In seiner Pracht konnte ich nicht nehmen, da es sich auf *Nacht* reimt.

5. a) Die Soldaten lieferten sich eine erbitterte Schlacht. Auch Mark gehörte zum Heer (13) und warf auf die Feinde _____. [den/einen Speer] (12)
- b) Die Soldaten lieferten sich eine erbitterte Schlacht. Auch Mark gehörte *zum Trupp* (15) und warf auf die Feinde _____.

Truppe, Gruppe, Armee, Kompanie und *Bataillon* habe ich nicht genommen, da diese Begriffe oft Reimwörter besitzen (*Suppe, Schnee, Ski, Spion* usw.), auch wenn diese Reimantworten semantisch ganz abwegig wären.

6. a) Lisa hatte es geschafft, einen sehr schönen Brief zu verfassen. Dies lag jedoch nicht an dem Stift (13), sondern an _____. [der/ihrer Schrift] (12)
- b) Lisa hatte es geschafft, einen sehr schönen Brief zu verfassen. Dies lag jedoch nicht an *dem Text* (10), sondern an _____.

Feder und *Papier* konnte ich nicht nehmen, da sich *Leder* und *ihr* darauf reimt.

7. a) Im ganzen Land hatte ein schreckliches Unwetter geherrscht und fast alle Gebäude wurden zerstört. Nur nicht der Turm (11), denn der widerstand _____. [dem Sturm] (10)

- b) Im ganzen Land hatte ein schreckliches Unwetter geherrscht und fast alle Gebäude wurden zerstört. Nur nicht *der Tempel* (11), denn der widerstand ____.

Ablenker mit einem femininen Genus wie beispielsweise *Hütte, Brücke, Kirche, Warte, Festung* usw. wollte ich nicht nehmen, da ich dann den Satz auf *denn **die** widerstand* ____ hätte ändern müssen.

8. a) Hanna und Erich hatten bereits seit langer Zeit gespielt. Nun hatten sie genug von dem Schach (13) und legten es zurück in ____ [das Fach] (11)
b) Hanna und Erich hatten bereits seit langer Zeit gespielt. Nun hatten sie genug von *dem Spiel* (15) und legten es zurück in ____.

Mühle und *Dame* habe ich nicht genommen, da es ungewöhnlich klingt, dass *sie genug von Dame/Mühle* haben.

9. a) Maria pflückte ein paar Blätter und Pflanzen für ihre Freundin. Die war nämlich eine Fee (15) und aß deswegen gerne ____ [Klee] (14)
b) Maria pflückte ein paar Blätter und Pflanzen für ihre Freundin. Die war nämlich *eine Elfe* (17) und aß deswegen gerne ____.
10. a) Tom hatte am Abend einen Mordshunger. Zuerst bestellte er Huhn mit Reis (12) und als Nachspeise gab es ____ [Eis] (10)
b) Tom hatte am Abend einen Mordshunger. Zuerst bestellte er Huhn mit *Gemüse* (11) und als Nachspeise gab es ____.

Nudeln konnte ich nicht nehmen, da es sich auf *Strudeln* reimt.

Hier noch einmal die Prime-Wörter und die mögliche Reimantwort der einsilbigen Wörter im Überblick:

	Reim-Prime	Reimantwort des Probanden	Nichtreim-Prime
1	Wand	Hand	Vorhang
2	Wicht	Licht	Gnom
3	Lohn	Sohn	Honorar
4	Glanz	Tanz	Schönheit
5	Heer	Speer	Trupp
6	Stift	Schrift	Text
7	Turm	Sturm	Tempel
8	Schach	Fach	Spiel
9	Fee	Klee	Elfe
10	Reis	Eis	Gemüse

Tabelle 1: Items mit einsilbigen Reimen

A.2. Items mit zweisilbigen Reimen

Die Itemanzahl für Satzpaare mit zweisilbigen Reimen soll ebenfalls 10 betragen ($n = 10$).

1. a) Der Kater Idefix jagte gerne kleine Nagetiere. Gestern fing er eine Ratte (15) und legte sie vor die Tür auf _____. [die Matte] (14)
- b) Der Kater Idefix jagte gerne kleine Nagetiere. Gestern fing er *eine Maus* (12) und legte sie vor die Tür auf _____.
2. a) Der ausländische Diplomat wurde zur Gala abgeholt. Er richtete sich noch schnell den Kragen (13) und stieg dann in _____. [den Wagen] (9)
- b) Der ausländische Diplomat wurde zur Gala abgeholt. Er richtete sich noch schnell *die Krawatte* (13) und stieg dann in _____.
3. a) Gustav machte Urlaub in einem fernen Land. Der dortige Arzt verabreichte ihm eine Spritze (14), denn er war kollabiert aufgrund _____. [der Hitze] (10)

- b) Gustav machte Urlaub in einem fernen Land. Der dortige Arzt verabreichte ihm *ein Medikament* (11), denn er war kollabiert aufgrund ____.

Tablette konnte ich nicht nehmen, denn es reimt sich auf *Skelette*.

4. a) Gerda war nach einem langen Marsch durstig geworden. Sie nahm einen Schluck aus der Flasche (11) und legte sie zurück in _____. [die Tasche] (10)
- b) Gerda war nach einem langen Marsch durstig geworden. Sie nahm einen Schluck aus *der Karaffe* (17) und legte sie zurück in _____.

Dose, Tasse konnte ich nicht nehmen, denn sie reimen sich auf *Hose, Kasse*. *Behälter* wollte ich nicht nehmen, da man dadurch den zweiten Teil zu *und steckte **ihn** zurück in _____*. ändern müsste.

5. a) Zum Geburtstag bekam Verena viele schöne Geschenke. Sie spielte nun jeden Tag mit ihrer Puppe (13) und fütterte sie mit _____. [Suppe] (13)
- b) Zum Geburtstag bekam Verena viele schöne Geschenke. Sie spielte nun jeden Tag mit *ihrer Barbie* (15) und fütterte sie mit _____.
6. a) Der Stammeshäuptling kämpfte mit seinem Feind. Dieser bekam eine Beule (16), denn der Häuptling traf ihn mit _____. [der Keule] (15)
- b) Der Stammeshäuptling kämpfte mit seinem Feind. Dieser bekam *eine Wunde* (13), denn der Häuptling traf ihn mit _____.
7. a) Das Konzert begann und die Fans warteten gespannt auf den Star. Der Bühnenarbeiter betätigte einen Hebel (13) und die Bühne wurde erfüllt von _____. [Nebel] (12)
- b) Das Konzert begann und die Fans warteten gespannt auf den Star. Der Bühnenarbeiter betätigte *einen Knopf* (13) und die Bühne wurde erfüllt von _____.

Schalter konnte ich nicht nehmen, da es sich auf *Falter* reimt.

8. a) Der kleine Ralf wünschte sich nichts sehnlicher als ein Haustier. Eines Tages bekam er einen Kater (13) und bedankte sich herzlich bei _____. [seinem Vater] (8)

- b) Der kleine Ralf wünschte sich nichts sehnlicher als ein Haustier. Eines Tages bekam er *einen Hamster* (15) und bedankte sich herzlich bei ____.
- Katze* konnte ich nicht nehmen, da es sich auf *Matze* (Spitzname für Matthias) reimt.
9. a) Die Familie bereitete sich auf das Weihnachtsfest vor. Der Vater schmückte die Tanne (15) und die Mutter briet Würstchen in _____. [der Pfanne] (14)
- b) Die Familie bereitete sich auf das Weihnachtsfest vor. Der Vater schmückte *den Baum* (10) und die Mutter briet Würstchen in _____.
10. a) Dem Söldner wurde großes Unrecht getan und er landete im Kerker. Jetzt wollte er Rache (12) und tötete mit einer List _____. [die/seine Wache] (12)
- b) Dem Söldner wurde großes Unrecht getan und er landete im Kerker. Jetzt wollte er *Vergeltung* (13) und tötete mit einer List _____.

Hier die Prime-Wörter und die mögliche Reimantwort der zweisilbigen Wörter im Überblick:

	Reim-Prime	Reimantwort des Probanden	Nichtreim-Prime
1	Ratte	Matte	Maus
2	Kragen	Wagen	Krawatte
3	Spritze	Hitze	Medikament
4	Flasche	Tasche	Karaffe
5	Puppe	Suppe	Barbie
6	Beule	Keule	Wunde
7	Hebel	Nebel	Knopf
8	Kater	Vater	Hamster
9	Tanne	Pfanne	Baum
10	Rache	Wache	Vergeltung

Tabelle 2: Items mit zweisilbigen Reimen

A.3. Füllitems

Nun folgen die 40 Füllitems, die an die eigentlichen experimentellen Sätze oben angepasst wurden. Die Primes in diesen Sätzen wurden wieder kursiv markiert.

Die Sätze, die mit den Items mit einsilbigen Reimantworten gematcht wurden:

1. a) Bernhard ging am Freitag auf ein Familienfest. Die Stimmung war gut, alle tranken *Wein* und lachten über ____.
- b) Am Donnerstag fand eine Halloweenparty statt. Es war gruselig, Eva war *eine Hexe* und ihre Freundin ____.
2. a) In seiner Kindheit hatte Gabriel einmal einen Alptraum. Er träumte Schreckliches *in jener Nacht* und schrieb es in ____.
- b) Agathe wollte heute unbedingt mit dem Zug abreisen. Doch vorm Bahnhof befand sich *eine Bank* und auf der saß ____.
3. a) Frank war den anderen Läufern weit überlegen. Bald sah er *das Ziel* und gewann somit auch dieses Jahr ____.
- b) Ruth bekam einen Geschenkskorb mit ein paar Früchten. Sie schälte *die Orange* und gab ihrer jüngeren Schwester ____.
4. a) Von Weitem schien das Bürogebäude vollkommen intakt zu sein. Es gab allerdings Mängel *beim Dach* und darum kümmerte sich nun schon ____.
- b) Der Geschäftsmann wollte wieder einmal seine Macht demonstrieren. Er arrangierte daher *ein Treffen* und die Gäste waren allesamt ____.
5. a) Die Kinder spielten im Winter gerne draußen im Freien. Sie bewarfen sich mit *Schnee* und fuhren auch oft mit ____.
- b) Der Theaterverein warb für seine neue Aufführung. Auch Flo war Mitglied *im Club* und gab den Passanten ____.
6. a) Tina glaubte zu wissen, wo genau sich ihr Bruder versteckt hielt. Dieser war aber nicht *im Bad*, sondern in ____.
- b) Hubert wollte unbedingt eine richtige Rakete basteln. Seine flog aber nicht *ins All*, sondern in ____.

7. a) Das ganze Dorf half mit bei den Vorbereitungen für das Fest und die Gerichte schmeckten wundervoll. Nur nicht *der Brei*, denn der roch stark nach ____.
- b) Der Windsturm von letzter Woche hatte schlimme Folgen und viele Bauern verloren ihre Ernte. Nur nicht *Raphael*, denn sein Feld lag bei ____.
8. a) Laura und Werner unterhielten sich über Sportarten. Werner spielte mittwochs immer *Golf* und Laura machte gerne ____.
- b) Tina und Simon trafen sich fast jeden Nachmittag. Sie tranken zusammen *einen Tee* und spielten eine Partie ____.
9. a) Der Jüngling war die einzige Rettung für das verzauberte Schloss. Er verspürte *keine Angst* und vertrieb letztendlich auch ____.
- b) Beim Entrümpeln des alten Schrankes kam eine Kostbarkeit zum Vorschein. Susi öffnete *die Dose* und entdeckte überrascht ____.
10. a) Fritz hatte einen gemütlichen Abend. Zuerst brät er sich Eier mit *Speck* und danach besuchte ihn ____.
- b) Am Abend kam Vera müde nach Hause. Nun freute sie sich auf *ein Stück Pastete* und genoss die Speise in ____.

Die Sätze, die mit den Items mit zweisilbigen Reimantworten gematcht wurden:

1. a) Der berühmte Zauberer Merlin braute einen Liebestrank. Er benötigte noch *Ingwer* und den besorgte er sich bei ____.
- b) Trotz schwerer Verluste gaben die fremden Soldaten nicht auf. Sie kämpften noch *einen Tag* und gewannen dann schlussendlich ____.
2. a) Der gefährliche Bankräuber konnte von der Polizei noch nicht geschnappt werden. Beim Überfall trug er *eine Maske* und er erschoss ____.
- b) Der Sportler fuhr übers Wochenende in ein Wellnesshotel. Zunächst gönnte er sich *eine Massage* und ging dann in ____.
3. a) Die Klein-Brüder waren erfolgreiche Männer. Im April gründeten sie in London *eine Firma*, denn sie erhofften sich dadurch ____.

- b) Kathrin hatte endlich genug Geld beisammen. Sie zog in die Hauptstadt und studierte dort *Biologie*, denn sie war interessiert an ____.
4. a) Der neue Sheriff war viel beliebter als der alte. Er hielt sich sehr streng an *das Gesetz* und befreite die Stadt von ____.
- b) Jans neuer Fall war eine große Herausforderung. Er war bereits seit Jahren *Detektiv* und nun auf der Suche nach ____.
5. a) Leonie wollte in ihrem Zimmer einen Aufsatz schreiben. Doch sie hatte gar *keine Idee* und spielte lieber mit ____.
- b) Dominik glaubte, seine Bestimmung gefunden zu haben. Er lebte nun schon seit langer Zeit *im Kloster* und war zuständig für ____.
6. a) Paul hatte sich den Abend frei genommen. Er fuhr mit dem Bus *ins Kino*, denn ihn interessierte ____.
- b) Kurt hatte Laura um Hilfe gebeten. Sie buk für ihn *einen Kuchen*, denn am Samstag ging er auf ____.
7. a) Die Vorstellung fing an und die Schauspieler waren sehr nervös. Einer der Hauptdarsteller verpasste *seinen Einsatz* und verließ daraufhin zornig ____.
- b) Die Neueröffnung des Restaurants war ein wirklicher Erfolg. Manche Gäste standen sogar draußen vor *der Tür* und warteten sehr geduldig auf ____.
8. a) Tobias wurde nachts vom Klingeln des Telefons geweckt. Claudia hatte leider *einen Unfall* und die Rettung brachte sie in ____.
- b) Herr Huber war schon immer ein Frühaufsteher gewesen. Im Morgengrauen ging er schon *zur Arbeit* und kaufte sich unterwegs noch ____.
9. a) Lili und ihre Schwester besuchten entfernte Verwandte. Diese besaßen *ein Hotel* und die Mädchen halfen mit bei ____.
- b) Onkel Leonard feierte seinen sechzigsten Geburtstag. Er öffnete *ein Fass Bier* und die Musikgruppe spielte ____.
10. a) Silvia hoffte sehr auf ihr Glück und machte bei einem Gewinnspiel mit. Sie gewann *die Reise* und flog mit ihren Eltern nach ____.

- b) Als Rentner bekam Wolfgang vieles verbilligt und nutzte diesen Vorteil.
Er ging oft *ins Museum* und betrachtete aufmerksam ____.

A.4. Übungsisems

1. Frau Weber lief jeden Tag eine Runde im Park. Eines Morgens fand sie eine *Jacke* und brachte sie zu ____.
2. Robert wollte sein Auto verkaufen und es meldete sich auch bald eine Interessentin. Er machte ihr ein *Angebot* und hoffte auf ____.
3. Stefanie sollte das Tochterunternehmen in Hamburg besuchen. Doch sie steckte leider im *Stau* und benachrichtigte ____.

A.5. Beschreibung

Folgender Text wurde den Probanden im Online-Fragebogen als Beschreibung dargeboten, bevor zu den Übungsisems übergegangen wird:

„Ab der nächsten Seite werden Sie nun immer zwei Sätze sehen, die jeweils ein kurzes Szenario beschreiben. Am Ende des zweiten Satzes befindet sich eine Lücke. Bitte lesen Sie die Sätze aufmerksam durch und tippen Sie in die Lücke, was den Satz Ihrer Meinung nach am besten vervollständigt. Dabei soll Ihre Antwort immer **ein Hauptwort** beinhalten.

Nehmen Sie ruhig die erste Antwort, die Ihnen einfällt!

Ein Beispiel hierfür wäre:

Markus ging am Samstag auf ein Konzert. Doch er war sehr in Eile und auf dem Hinweg verlor er ____.

Mögliche Eingaben Ihrerseits könnten lauten:

den Schlüssel
einen Ausweis

sein Ticket

Geld

usw.

Bitte benutzen Sie, wenn möglich, **keine** zusammengesetzten Wörter, wie: die Eintrittskarte (Eintritts-karte)

Es folgen drei Übungsitems und danach die Items des Experiments. Viel Vergnügen!

Wenn Sie erfahren möchten, was das genaue Thema der Untersuchung ist, können Sie am Ende des Fragebogens Ihre E-Mail-Adresse hinterlassen.“

B. Zusammenfassung

In dieser Arbeit werden Reime im Zusammenhang mit ihrer Produktion bzw. mit dem Zugriff auf das Lexikon behandelt. Zunächst wird der phonologische Aufbau von Reimen und Silben erläutert und gezeigt, wie ein Reim aufgebaut sein kann. Es werden Theorien mehrsilbiger Reime vorgestellt. Danach folgt ein Überblick über aktuelle Forschungsergebnisse mit Reimen bei Kindern, sprachgesunden und aphasischen Erwachsenen, sowie ein Einblick in die neurologische Verarbeitung von Reimen. Themen wie die phonologische Bewusstheit und der Schriftspracherwerb werden angeschnitten. Im vierten Kapitel wird auf das mentale Lexikon und auf den Zugriff darauf eingegangen. Serielle Sprachproduktionsmodelle werden interaktiven Sprachproduktionsmodellen gegenübergestellt. Schließlich folgt das Experiment, das in Anlehnung an Rapp und Samuel (2002) im Rahmen dieser Arbeit durchgeführt wurde, um die Produktion von Reimen in Sätzen zu untersuchen. Die Ergebnisse zeigen keinen Effekt der Phonologie auf die Wortwahl, mögliche Ursachen dafür werden besprochen.

C. Lebenslauf

PERSÖNLICHE DATEN

Name:	Michaela Rausch
Geburtsdatum:	01.06.1988
Geburtsort:	Kirchdorf an der Krems, Oberösterreich
Staatsangehörigkeit:	Österreich
E-Mail:	michaela.rausch@gmx.at

AUSBILDUNG

1994 - 1998	Volksschule Wartberg/Krems
1998 - 2006	Gymnasium der Abtei Schlierbach
seit Oktober 2006	Universität Wien, Österreich Studium der Sprachwissenschaft (Diplomstudium) Schwerpunkt: Psycho-, Patho- und Neurolinguistik
Oktober 2008 - Juli 2009	Universität Potsdam, Deutschland Studium der Patholinguistik (Bachelorstudium) ERASMUS-Auslandsaufenthalt
September 2011 - Juni 2012	Balassi Intézet, Budapest, Ungarn Studium der Hungarologie (Stipendium) verifiziert von der Universität Pécs
seit Oktober 2012	Ludwig-Maximilians-Universität München, Deutschland Studium der Computerlinguistik (Bachelorstudium) Nebenfach: Sprache, Literatur, Kultur

BERUFLICHE ERFAHRUNGEN (AUSWAHL)

März/April 2009	dreiwöchiges Praktikum am Berlin NeuroImaging Center (BNIC) Charité, Berlin, Deutschland
Mai - Mitte Juli 2009	wissenschaftliche Hilfskraft am Berlin NeuroImaging Center (BNIC) Charité, Berlin, Deutschland
Mai 2009	zweitägiges Praktikum auf der HNO-Station im Ernst von Bergmann Klinikum Potsdam, Deutschland
August 2009	einmonatiges Praktikum am Institut für Sinnes- und Sprachneurologie Konventhospital Barmherzige Brüder Linz, Oberösterreich
seit Dezember 2009	Mitarbeit beim ZIT-Projekt: „Gegen die Sprachlosigkeit mit linguistisch fundierten Sprachtest- und Sprachtrainingsverfahren (ELA Photo Series)“ Konzipierung eines Computerprogramms zur Sprachtherapie Projektleitung: Dr. Jacqueline-Ann Stark
seit August 2012	Werkvertrag bei der österreichischen Akademie der Wissenschaften (ÖAW)

SPRACHKENNTNISSE

Deutsch	Muttersprache
Englisch	8-jährige Schulausbildung
Ungarisch	Kurse an der Universität Wien, Auslandsjahr, Stufe C1 nach dem Gemeinsamen Europäischen Referenzrahmen für Sprachen
Französisch	4-jährige Schulausbildung, Aufbaukurs an der Universität Wien
Spanisch	Grundkurs an der Universität Wien
Latein	6-jährige Schulausbildung