



universität
wien

Diplomarbeit

Titel der Diplomarbeit

Sublexikales Wortlesen: Eine Analyse der Verarbeitung von Onset/Rime und Konsonantencluster in der 2. und 4. Schulstufe

Verfasserin

Julia Kast

Angestrebter akademischer Grad

Magistra der Naturwissenschaften (Mag. rer. nat.)

Wien, 2013

Studienkennzahl: 298

Studienrichtung: Psychologie

Betreuer: Ao. Univ.-Prof. Mag. Dr. Alfred Schabmann

Danksagung

Zu Beginn möchte ich mich bei einigen Personen bedanken, die mir bei der Erstellung dieser Diplomarbeit geholfen und unterstützt haben.

Allen voran danke ich meinem Betreuer Ao. Univ.-Prof. Mag. Dr. Alfred Schabmann für die unkomplizierte und produktive Zusammenarbeit.

Besonders bedanken möchte ich mich auch bei meiner Familie und meinem Freund für die Unterstützung während des Studiums.

Großer Dank gebührt auch den Lehrkräften, SchülerInnen und Eltern der teilnehmenden Volksschulen. Besonders möchte ich mich bei den jeweiligen Direktorinnen bedanken, welche offen und freundlich auf mein „Vorhaben Diplomarbeit“ reagiert und mir ein problemloses Forschen ermöglicht haben:

Gertrude STANZL, VS Brehmstraße Wien

Monika HAUTZINGER, VS Mönchhof

Petra SCHMIDT, VS Zurndorf

Emilie WAHRMANN, VS Gols

Inhaltsverzeichnis

1. Einleitung.....	7
2. Theoretischer Hintergrund.....	8
2.1 Was versteht man unter sublexikalen Clustern?.....	8
2.2 Modelle des Leseerwerbs.....	11
2.3 Theorien des kompetenten Lesens.....	16
2.3 Bisherige Studien zur Verwendung sublexikaler Cluster.....	22
3. Untersuchung.....	31
3.1. TeilnehmerInnen.....	31
3.2. Instrumente.....	31
2.3. Durchführung.....	33
4. Ergebnisse.....	35
5. Diskussion.....	42
5.1 Interpretation.....	42
5.2 Kritische Diskussion.....	48
5.3 Ausblick und praktische Relevanz.....	50
6. Zusammenfassung.....	51
6.1 Deutsche Zusammenfassung.....	51
6.2 Summary.....	53
7. Literaturverzeichnis.....	55
8. Anhang.....	57
Anhang A.....	57
Anhang B.....	58

1. Einleitung

Lesen zählt zu einer der wichtigsten Grundkompetenzen zur Bewältigung des Alltags. Aus diesem Grund ist es ein wesentliches Bildungsziel elementarer Bildungseinrichtungen und darüber hinaus in den meisten Unterrichtsgegenständen fest verankert. Dem Leseverständnis, dem sinnerfassenden Aspekt dieser Fertigkeit, geht die Entwicklung der Lesekompetenz voraus.

Die OECD (2005) beschreibt die Lesekompetenz als eine Fähigkeit, schriftliches Material zu verstehen, zu nutzen und darüber zu reflektieren. Diese soll nicht nur dazu dienen, das eigene Wissen und Potential zu erweitern, sondern auch dazu befähigen, am gesellschaftlichen Leben teilzunehmen.

Bevor sich ein sinnverstehendes Lesen entwickeln kann, müssen basale Lesefähigkeiten angeeignet werden. Die Gestaltung des Erstleseunterrichts stellt hierbei eine wesentliche Komponente dar. Ein Verständnis für die verschiedenen Theorien des Leseerwerbs und für die theoretischen Aspekte, wie Wörter im Zuge des Leseprozesses verarbeitet werden, ist dabei unabdingbar.

Diese Arbeit soll einen Beitrag zu bisherigen Forschungsbefunden leisten, inwieweit Onset und Rime (An- und Auslaut einer Silbe) und Konsonantencluster als Einheiten beim lauten Lesen verarbeitet werden.

2. Theoretischer Hintergrund

2.1 Was versteht man unter sublexikalen Clustern?

Marinus und de Jong (2008) beschreiben sublexikale Cluster als Einheiten, welche größer als ein Buchstabe, jedoch kleiner als ein Wort sind. Zu dieser Gruppe gehören unter anderem Konsonantencluster („str“), ein Digraph (zwei Buchstaben, die für ein Phonem oder eine Silbe stehen-wie zum Beispiel „ch“), Silben („lu“ in Lupe), Onsets (Anlaut einer Silbe-„s“ in Seil) oder Rimes (der verbleibende Rest oder Auslaut der Silbe-„eil“ in Seil).

Theorien des Leseerwerbs haben unterschiedliche Auffassungen darüber, wie diese sublexikalen Einheiten repräsentiert und verarbeitet werden. Entwicklungsmodelle, wie jenes von Ehri (1995), gehen davon aus, dass die Verarbeitung und der Gebrauch sublexikaler Cluster erlernt werden und dass diese Verknüpfungen im mentalen Gedächtnis fest etabliert sind. Im Dual-Route Modell (Coltheart, Rastle, Perry, Langdon & Ziegler, 2001) werden Wörter entweder direkt aus einem mentalen Lexikon abgerufen oder Buchstabe für Buchstabe durch die phonologische Rekodierung verarbeitet. Sublexikale Cluster werden dabei zwar bei der phonologischen Repräsentation berücksichtigt, jedoch nicht im Modell gespeichert. Netzwerkmodelle, wie jenes von Seidenberg und McClelland (1989), gehen wiederum von einem einzigen Zugangsweg aus, bei dem häufig auftretende Verbindungen von Buchstaben im Modell gespeichert werden. Sublexikale Cluster werden demnach mit der Zeit erlernt und beim Leseprozess berücksichtigt.

Die in den Modellen beschriebenen Verarbeitungsmöglichkeiten von sublexikalen Clustern wurden in verschiedenen Studien untersucht. Der Fokus lag hierbei unter anderem darauf, ob diese beim Lesen als eine Einheit wahrgenommen und als solche verwendet werden. Jener Sachverhalt wurde beispielsweise bei Levitt, Healy und Fendrich (1991), Bowey (1990, 1996) oder Marinus und de Jong (2008) untersucht.

Weiters beschäftigten sich unter anderem Frederiksen und Kroll (1976) und Martensen, Maris und Dijkstra (2003) mit der Verarbeitung von Konsonanten- und Vokalcluster im Leseprozess.

Ein wesentliches Problem der bisherigen Studien ist, dass diese meist in einem anderen Sprachraum durchgeführt wurden. Durch die unterschiedliche Struktur, Komplexität und auch

Konsistenz der verschiedenen Sprachen, können demnach nur schwer vorbehaltlos Rückschlüsse in das Deutsche übertragen werden.

In dieser Studie soll der Schwerpunkt einerseits auf der Verarbeitung von Onset und Rime und andererseits auf jener von Konsonantencluster beim lauten Lesen liegen.

Es wird angenommen, dass Wörter beim Lesen silbenweise zergliedert werden. Dies wird dadurch unterstützt, dass phonologische Fertigkeiten durch das allseits bekannte „Silbenklatschen“ geschult und erlernt werden. Ein bedeutsamer Aspekt für das Lesen lernen ist es nun, herauszufinden, inwieweit kleinere Einheiten der Silben, wie Onset, Rime oder Konsonantencluster, ebenfalls als Einheiten wahrgenommen werden und die Leseflüssigkeit in weiterer Folge beeinflussen.

Nach Klicpera, Schabmann und Gasteiger-Klicpera (2010) spielt hierbei die phonologische Bewusstheit der Kinder eine wesentliche Rolle. Darunter versteht man die Fähigkeit, die Einzelsegmente der Sprache, wie beispielsweise Silben oder Phoneme (Laute), zu erfassen und auch wahrzunehmen. Nach Goswami (2000, zitiert nach Klicpera et al., 2010, S. 26) ist es für SchülerInnen erst nach Beginn des Erstleseunterrichts möglich, Wörter in ihre Phoneme zu zerlegen. Vor dieser Zeit fällt es den Kindern leichter, Wörter silbenweise zu segmentieren. Dabei scheint es am einfachsten, den Silbenanfang vom verbleibenden Rest zu trennen.

Im weiteren Zusammenhang ist die Gestaltung des Erstleseunterrichts ein wesentlicher Faktor des Leseerwerbs. Nach Klicpera et al. (2010) ist es im deutschsprachigen Raum eine gängige Unterrichtsmethode, zuerst einzelne Buchstaben zu erlernen, womit in weiterer Folge Wörter beziehungsweise Sätze gelesen werden können. Da den SchülerInnen durch diese sogenannten Graphem-Phonem-Korrespondenzen beigebracht wird, dass Buchstaben als Repräsentationen von Lauten dienen, wächst somit eine tiefgründige Einsicht in die phonologische Struktur der Sprache.

Vor allem bei Sprachen, die eine unregelmäßige Graphem-Phonem-Korrespondenz aufweisen (wie dies beispielsweise im Englischen der Fall ist), scheint es für SchülerInnen hilfreich zu sein, eine Gliederung in An- und Auslaut des Wortes beim Erlernen zu verwenden. Im Unterricht können die Kinder demnach explizit darauf hingewiesen werden, auf den Auslaut

(den reimenden Teil eines Wortes) und dessen Ähnlichkeit mit anderen Wörtern zu achten. Goswami (1999) betont in diesem Sinne, dass es einen positiven Einfluss auf den Leseprozess haben kann, wenn Kinder dahingehend geschult werden, Wörter in Onset und Rime zu zerlegen, um diese phonologischen Sequenzen mit der Schreibweise in Verbindung zu bringen. LeseanfängerInnen haben demnach diese kleinsten phonologischen Cluster gespeichert, wodurch das Lesen erleichtert wird. SchülerInnen mit einer verminderten Lesefertigkeit könnte dieser Ansicht nach durch konkrete Instruktionen und der Betonung dieser kleineren sprachlichen Einheiten geholfen werden.

Mit entsprechenden Förderprogrammen kann schon frühzeitig verhindert und auch dagegen vorgebeugt werden, dass eine Beeinträchtigung der Lesefähigkeit auftritt. Eines der wenigen Trainings, welches sich unter anderem auch mit der Verarbeitung von Onset und Rime beschäftigt, nennt sich PHAB/DI (Phonological Analysis and Blending/Direct Instruction-Phonologische Analyse und Zusammenlauten/Direkter Unterricht; Lovett 1999, zitiert nach Klicpera et al., 2010, S. 250). Durch eine Teilnahme an diesem Training können nach Klicpera et al. (2010) erhebliche Fortschritte erzielt werden. Diese sind vor allem bei den phonologischen Fertigkeiten, wie den Graphem-Phonem-Zuordnungen, zu beobachten. Eine bemerkenswerte Leistungssteigerung hinsichtlich der Verarbeitung von Onset und Rime eines Wortes wird von den Autoren/der Autorin nicht erwähnt.

Ein sogenanntes Rime-Training wurde von Wise, Olson und Treiman (1990, zitiert nach Goswami, 1999, S. 229) verfasst. Obwohl die Ergebnisse nur über einen kürzeren Zeitraum beobachtet werden konnten, ließ sich feststellen, dass das Wortlernen der TeilnehmerInnen bessere Resultate lieferte, wenn eine Teilung zwischen Onset und Rime vorgenommen wurde, als wenn diese im Rime nach dem Vokal eingeführt wurde.

Ein weiteres nennenswertes Förderprogramm wurde von Reuter-Liehr (2001, zitiert nach Klicpera et al., 2010, S. 253) erstellt. Der Fokus liegt hierbei auf dem silbenweisen Sprechen und Lesen. Dabei sollen Bewegungen des ganzen Körpers, wie Arm- und Schrittbewegungen, beim silbenweisen Aneignen von Wörtern integriert werden. Die unterschiedlichen Phonemstufen des Förderprogramms, welche einen steigenden Schwierigkeitsgrad besitzen, haben nicht nur die Erarbeitung von An- und Auslaut, sondern auch die der häufigen Konsonantencluster zum Inhalt.

In der vorliegenden Arbeit soll vorerst ein kurzer Überblick über die verschiedenen Entwicklungsmodelle des Lesens gegeben werden. Danach folgt eine kurze Darstellung von Theorien des kompetenten Lesers/der kompetenten Leserin, welche die Worterkennung und dessen Verarbeitung beschreiben und in ihren Grundannahmen teilweise verschieden sind. In den beiden Aspekten soll in jeder Hinsicht der Fokus auf der Verarbeitung und der Verwendung sublexikaler Cluster liegen. Abschließend folgt eine beispielhafte Auflistung von Studien, die sublexikale Einheiten in die Forschung integriert haben.

2.2 Modelle des Leseerwerbs

Theorie des Leseerwerbs (Frith, 1985)

Friths Annahmen stellen eine Weiterentwicklung der kognitiv-entwicklungspsychologischen Theorie von Marsh, Friedman, Welch und Desberg (1981, zitiert nach Frith 1985, S. 305) dar. Frith geht davon aus, dass sich die Entwicklung des Lesens in drei Phasen gliedert, welche wiederum durch drei verschiedenartige Strategien gekennzeichnet sind.

In der *logografischen Phase* können bekannte Wörter durch visuelle Merkmale rasch erkannt werden. Dabei dient der Anfangsbuchstabe oder die Schriftart als ein wesentlicher Hinweis. Die Abfolge der Buchstaben und die phonologischen Aspekte des Wortes sind dabei zweitrangig.

In der *alphabetischen Phase* wissen die Kinder bereits über Graphem-Phonem-Zuordnungen Bescheid und können diese auch anwenden. Wörter werden hierbei systematisch analysiert, indem sie Buchstabe für Buchstabe dekodiert werden. Die Buchstabenabfolge und die phonologischen Faktoren spielen hierbei eine wesentliche Rolle. In der alphabetischen Phase, welche meist zu Schuleintritt erreicht wird, können nicht nur neue, sondern auch Pseudowörter erlesen werden.

Die *orthografische Phase* ermöglicht es dem Leser/der Leserin, Wörter rasch in orthografische Einheiten zu zerlegen, ohne eine phonologische Rekodierung anzuwenden. Dadurch können ganze Wörter oder Morpheme (lautliche Einheiten mit einer Bedeutung) direkt erkannt und identifiziert werden. Eine buchstabenweise Zerlegung und eine darauffolgende Rekodierung sind in dieser Phase nicht mehr notwendig.

Friths Annahmen lassen durchaus auf eine Abspeicherung und auf eine Verwendung sublexikaler Cluster beim Lesen schließen.

Phasentheorie (Ehri, 1995)

Ehri (1995) unterscheidet vier verschiedene Phasen des Leseerwerbs, die in Abhängigkeit des jeweiligen Stadiums mit dem alphabethischen System in Verbindung stehen. Jede dieser Phasen besitzt eine bestimmte Übertragung zwischen dem geschriebenen Wort und dessen Aussprache und der damit verknüpften Bedeutung im Gedächtnis.

Präalphabethische Phase: Hierbei werden Wörter erkannt, indem, ähnlich wie bei Frith (1985), von visuellen Merkmalen auf die Aussprache beziehungsweise die Bedeutung rückgeschlossen wird. Diese Aspekte werden im Gedächtnis gespeichert und können bei Bedarf wieder abgerufen werden. In dieser Phase wird zum Beispiel der Schriftzug „McDonalds“ nicht mithilfe der Buchstaben erlesen, sondern durch den großen goldenen Bogen identifiziert und erkannt. Auch wenn Buchstaben im Wort verändert werden, würde dies nicht bemerkt werden. In dieser präalphabethischen Phase hat sich demnach noch kein Verständnis über eine Buchstaben-Laut-Zuordnung entwickelt.

Partiell alphabethische Phase: Leseanfänger können in dieser Phase teilweise alphabethische Verbindungen zwischen einigen gelesenen Buchstaben und deren Aussprache herstellen. Hierbei sind vor allem der erste und der letzte Buchstabe hervorstechend, da dieser als wesentliches Wiedererkennungsmerkmal dient.

Vollalphabetische Phase: In der vollalphabetischen Phase werden Wörter erlesen, indem bereits vollständige Verbindungen der Buchstaben-Lautzuordnung hergestellt werden können. Die LeserInnen haben in dieser Phase bereits ein Verständnis dafür entwickelt, welche Buchstaben und welche Laute zusammengehören. Die Lesegenauigkeit kann durch diesen Vorgang verbessert werden. Darüber hinaus können Unsicherheiten bei ähnlich geschriebenen und auch bei unbekannten Wörtern vermieden werden, da solche immer wieder im Gedächtnis abgespeichert werden und sich somit das mentale Lexikon stetig erweitert.

Konsolidierte alphabetische Phase: Hierbei wird das Leseverhalten weiter automatisiert und Buchstabengruppen wie Silben, Onsets oder Rimes werden als Einheiten im Gedächtnis gespeichert, wodurch das Erlesen von Wörtern beschleunigt werden kann.

Ehri (1995) geht davon aus, dass sublexikale Cluster wie zum Beispiel Silben, Onsets oder auch Rimes als visuelle Einheiten im Gedächtnis gespeichert und bei Bedarf abgerufen werden.

Psycholinguistische Grain Size Theorie (Ziegler & Goswami, 2005)

Die Grain Size Theorie von Ziegler und Goswami (2005) befasst sich im Wesentlichen mit der Größe von Spracheinheiten („grain sizes“) verschiedener Sprachkulturen und inwieweit diese für Probleme bei der Aneignung von Lesestrategien verantwortlich sind. Dabei spielt die phonologische Bewusstheit eine wichtige Rolle. In diesem Zusammenhang sind LeseanfängerInnen mit den folgenden drei Problemen konfrontiert:

Verfügbarkeitsproblem: Nicht alle phonologischen Einheiten sind vor dem Lesen bewusst verfügbar. Um die Schreibweise einer eindeutigen Aussprache zuzuordnen sind weitere kognitive Abläufe notwendig.

Konsistenzproblem: Es gibt oftmals mehrere Möglichkeiten, orthografische Cluster auszusprechen und einige phonologische Einheiten können auf verschiedene Weise geschrieben werden. Diese Inkonsistenz führt zu einer Verlangsamung des Leseprozesses.

Granularitätsproblem: Wenn das phonologische System größere sprachliche Einheiten („grain sizes“) verarbeiten muss, müssen auch mehr orthographische Einheiten erlernt werden. Dies basiert auf der Tatsache, dass es mehr Wörter als Silben, mehr Silben als Rimes und so weiter gibt.

Die Lesefertigkeit ist nicht nur von der Bewältigung der drei Problembereiche, sondern auch von den Charakteristika der jeweiligen Sprache abhängig. Drei Faktoren sind laut Ziegler und Goswami (2005) für diese kulturellen Unterschiede verantwortlich: die Konsistenz der Buchstaben-Lautzuordnung, die Granularität der orthographischen und phonologischen Einheiten und die Lehrmethoden.

Eine wesentliche Beziehung sehen Ziegler und Goswami (2005) zwischen den ersten beiden oben angeführten Problemen bei LeseanfängerInnen: der Konsistenz und der Granularität. Kleinere sprachliche Einheiten sind tendenziell inkonsistenter als größere. Konsistente Sprachen (wie Spanisch oder Italienisch) in denen beispielsweise ein Buchstabe immer gleich ausgesprochen wird, sind darüber hinaus einfacher zu erlernen als inkonsistente (wie beispielsweise Englisch). Der generelle Aspekt der Inkonsistenz führt in weiterer Folge dazu, dass größere sprachliche Einheiten gebildet werden müssen, um eine richtige phonologische Rekodierung zu erstellen. Dadurch müssen mehrere Einheiten erlernt werden, was wiederum mehr Zeit in Anspruch nimmt. Im Deutschen, einer sehr konsistenten Sprache, sollten in dieser Hinsicht vermehrt kleinere Einheiten vermittelt werden, da dies zu einem schnelleren Lernerfolg führt.

Die Grain Size Theorie spricht im Wesentlichen für das Lesen in Form von sublexikalischen Clustern (in diesem Fall von „grain sizes“). Vor allem bei inkonsistenten Sprachen scheint die Abspeicherung von sprachlichen Einheiten wesentlich zu sein, da dies den Leseprozess zunehmend erleichtert und beschleunigt. Diese Tatsache ist zwar im Deutschen von nicht allzu großer Bedeutung, da es sich hierbei um eine relativ konsistente Sprache handelt, jedoch scheinen geübte LeserInnen sublexikalische Einheiten abzuspeichern, um den Leseprozess zu vereinfachen.

Kompetenzentwicklungsmodell (Klicpera, Schabmann & Gasteiger-Klicpera, 2010)

Grundlage dieses Modells, welches sich auf den deutschsprachigen Raum bezieht, ist nicht die Abfolge von Entwicklungsphasen, sondern die zu erwerbenden Lesekompetenzen während der Entwicklung. Aus diesem Grund wird auch von einem Kompetenzentwicklungsmodell gesprochen. Weiters wird hier von unterschiedlichen Entwicklungsverläufen ausgegangen, welche von den individuellen Lernvoraussetzungen der SchülerInnen und der Instruktion im Unterricht abhängig sind.

Ähnlich wie beim Dual-Route Modell von Coltheart et al. (2001), welches im Kapitel 2.3 näher beschrieben wird, wird bei diesem Modell von einer sogenannten Zwei-Wege-Theorie des Worterkennens des kompetenten Lesers/der kompetenten Leserin ausgegangen. Ein Wort wird demnach entweder direkt aus dem mentalen Lexikon abgerufen (lexikalischer Weg), oder Buchstabe für Buchstabe mittels phonologischer Rekodierung (nicht-lexikaler Weg) ermittelt (siehe auch: Dual Route Modell nach Coltheart et al., 2001).

Beide Zugriffsarten entwickeln sich laut dem Kompetenzentwicklungsmodell in Interaktion mit der Leseinstruktion im Unterricht und individuellen Fördermaßnahmen.

Wie bei Ehri (1995) wird bei diesem Modell von einer Vorstufe, der sogenannten *präalphabetischen Phase*, ausgegangen, wobei noch keine Kenntnisse über die konkrete Buchstaben-Laut-Zuordnung vorhanden sind und Wörter aufgrund von Merkmalen (goldener Bogen beim Schriftzug „McDonald’s“) erkannt und identifiziert werden.

Die darauf folgende *alphabetische Phase mit geringer Integration* wird von den Autoren/der Autorin als erste „echte“ Leselernphase bezeichnet. Die zum Lesen benötigten Fähigkeiten sind noch nicht vollständig ausgebildet und die einzelnen Teilprozesse bilden noch kein funktionierendes Gesamtsystem. Diese Phase wird vor allem beim Schuleintritt erkennbar. Zu dieser Zeit erlernen SchülerInnen einerseits das alphabetische Prinzip und andererseits jenes der phonologischen Rekodierung. Die regelmäßigen Graphem-Phonem-Beziehungen der deutschen Sprache führen rasch zu ersten Lern- und Leseerfolgen. Parallel zum phonologischen Rekodieren entwickelt sich die Fähigkeit, ganze Wörter aus dem mentalen Lexikon abzurufen.

In der *alphabetischen Phase mit voller Integration* wird der Leseprozess, sowohl beim lexikalen als auch beim nicht-lexikalen Vorgehen, zunehmend automatisiert. Dies bedeutet, dass SchülerInnen beim Lesen weniger Fehler unterlaufen und ihre Geschwindigkeit zudem gesteigert werden kann. Es scheint, dass die Zunahme der Lesegeschwindigkeit mit der Entwicklung des partiell-lexikalischen Lesens zusammenhängt. Dabei werden mental Einheiten gebildet, die, wie zum Beispiel häufig vorkommende Buchstabencluster („ck“ oder „ch“), beinhalten. Diese Vorgangsweise ist demnach eine Fortsetzung des buchstabenweisen Rekodierens.

Die alphabetische Phase mit voller Integration stellt den Übergang in die letzte Phase dar: der *automatisierten und konsolidierten Integration* von allen beteiligten Verarbeitungsprozessen.

Das Kompetenzentwicklungsmodell von Klicpera et al. (2010) spricht auch für die Verwendung von Buchstabencluster während des Lesens, was sich wiederum auf die Lesegeschwindigkeit auswirkt.

2.3 Theorien des kompetenten Lesens

Dual-Route Cascaded Modell (Coltheart, Rastle, Perry, Langdon & Ziegler, 2001)

Das Dual-Route Cascaded Modell (DRC Modell) beschreibt den Vorgang, wie Wörter und auch Pseudowörter erkannt und letztlich verbalisiert werden. Dieser kann im Wesentlichen über zwei Wege ablaufen: die lexikalische oder die nicht lexikalische Route. Wird der *lexikalische Weg* aktiviert, wird das gelesene Wort direkt aus einem mentalen Lexikon abgerufen, in welchem bereits erlesene Wörter gespeichert sind. Die *nicht lexikale Route* entspricht dem Vorgang einer phonologischen Rekodierung, bei der Wörter buchstabenweise in Laute transferiert werden. Beim lexikalischen Weg kann zusätzlich dazu das semantische System zugeschaltet werden. Nach Klicpera et al. (2010) sind die einzelnen Pfade voneinander abhängig, jedoch nicht immer gleichzeitig verfügbar. Der Zugriff auf das mentale Lexikon ist bei der lexikalischen Verarbeitung nicht erfolgreich, wenn es sich um ein unbekanntes Wort handelt. Dies können unter anderem Fremd-, Pseudo- oder noch nicht erlernte Wörter sein. In diesem Fall muss das Schriftbild über die nicht lexikale Route durch eine phonologische Rekodierung ermittelt werden.

Jeder der Pfade besteht aus mehreren Ebenen, die sich gegenseitig beeinflussen. Diese wiederum beinhalten Einheiten („units“), die die kleinsten symbolischen Teile des Modells darstellen. In der Ebene des orthographischen Lexikons stellen Wörter diese units dar. Sie interagieren entweder durch gegenseitige Aktivierung oder Hemmung miteinander. Für einen detaillierten Überblick über die einzelnen Ebenen und Einheiten siehe Abbildung 1.

Eine Besonderheit dieses Zwei-Wege-Modells ist sein stufenweiser („cascaded“) Ablauf. Dies bedeutet, dass eine Aktivierung von einer Einheit direkt zur nächsten weitergeleitet wird, ohne dass ein bestimmtes Kriterium erreicht sein muss.

Die oben angeführten Grundlagen des DRC Modells wurden in eine computerisierte Form übertragen. Dadurch kann der menschliche Lesevorgang sozusagen simuliert und Prozesse, die während der visuellen Worterkennung und dem lauten Lesen aktiv sind, analysiert werden. Das Computerprogramm soll ein möglichst genaues Abbild der menschlichen Vorgänge darstellen.

Der Ablauf des Leseprozesses gestaltet sich in diesem Modell auf folgende Weise: Als Input dient das geschriebene Wort. Im Zuge der visuellen Wortanalyse werden die Buchstaben und dessen Eigenschaften analysiert und gleichzeitig dazu wird die Einheit der

Buchstabenrepräsentationen aktiviert. Danach teilen sich die Wege: der nicht lexikalische

Weg transferiert Buchstaben in Laute, wobei Buchstabengruppen wie „st“ oder „sch“ zusammen als ein Phonem wahrgenommen werden. Dieser Prozess verläuft immer von links nach rechts und dabei werden die Positionen der Buchstaben zusätzlich berücksichtigt.

Beim lexikalischen Weg wird in der Einheit des orthografischen Lexikons das ganze Wort aus dem mentalen Speicher abgerufen. Hierbei werden alle Buchstaben gleichzeitig verarbeitet: jeder einzelne hat eine aktivierende Wirkung für jene Wörter, in denen sie vorkommen und eine hemmende für jene, in denen sie nicht vorkommen.

Der Abruf eines Wortes aus dem orthografischen Lexikon aktiviert in weiterer Folge die Einheit des phonologischen Lexikons.

Beim sinnverstehenden Lesen wird bei dieser Route darüber hinaus noch das semantische System aktiviert und dazwischengeschaltet.

Sowohl die phonologische Rekodierung, als auch der Abruf aus dem Lexikon führen zur Phonemeinheit, welche letztendes die Aussprache aktiviert und generiert.

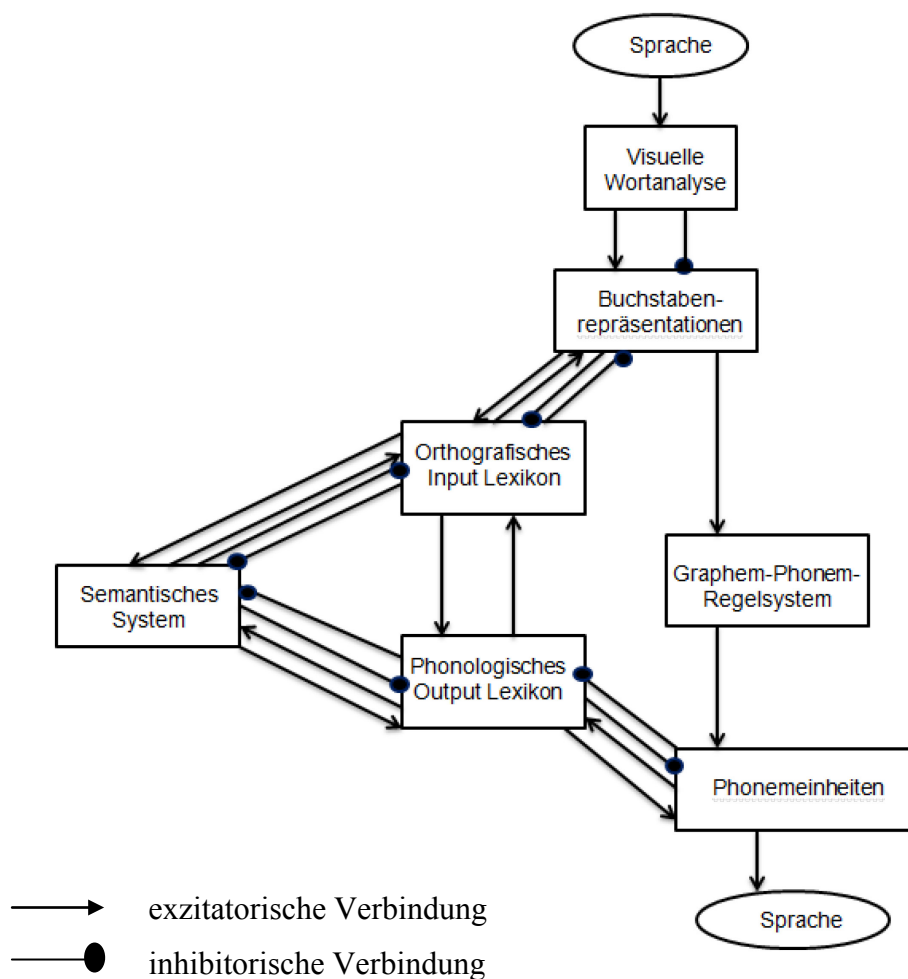


Abbildung 1: Das Dual-Route Cascaded Modell (Grafik adaptiert nach Coltheart et al., 2001).

Das DRC Modell kann einsilbige, jedoch keine mehrsilbigen Wörter, relativ gut verarbeiten. Sublexikale Cluster werden in diesem Modell zwar berücksichtigt, wenn die Buchstabengruppe einen Laut repräsentiert („sp“ oder „st“), jedoch werden diese nicht für eine weitere Verwendung als Einheit gespeichert. Eine Berücksichtigung sublexikaler Cluster während des Lesens ist nach dem DRC Modell demnach nicht eindeutig bestätigt.

Netzwerkmodelle

Parallel Distributed Modell (Seidenberg & McClelland, 1989; Plaut, McClelland, Seidenberg & Patterson, 1996)

Im Modell von Seidenberg und McClelland (1989) werden bei der Worterkennung im Wesentlichen drei Arten von Codes unterschieden: orthografische, phonologische und semantische. Dieses Netzwerk lässt sich demnach in orthografische, phonologische und semantische Repräsentationseinheiten einteilen, in denen unterschiedliche Informationen von Wörtern gespeichert sind. Der Worterkennungsprozess kann darüber hinaus von Kontextfaktoren, welche syntaktisch, semantisch oder pragmatisch bedingt sein können, beeinflusst werden. Die einzelnen Repräsentationen sind durch sogenannte verborgene Einheiten („hidden units“) miteinander verbunden. Dieses, sich daraus ergebende Netzwerk, ist in seinen Komponenten wechselseitig miteinander verknüpft und es besteht die Möglichkeit der gegenseitigen Beeinflussung.

Bei Netzwerkmodellen im Allgemeinen wird angenommen, dass sich der zuvor beschriebene Aufbau schrittweise durch die Auseinandersetzung des Lesers/der Leserin mit der Schriftsprache entwickelt. Dabei können verschiedenartige Informationen in den einzelnen Repräsentationseinheiten gespeichert werden, was wiederum einen schnellen und unbewussten Zugriff ermöglicht. Eine bildliche Darstellung der Erläuterungen kann Abbildung 2 entnommen werden. Der fett gedruckte Teil der Abbildung links stellt den Kern des Modells dar – dem sogenannten Triangelmodell, welches rechts gesondert dargestellt wird. Dabei werden die semantische und die kontextspezifische Einheit entfernt und somit besteht dieses Modell nur aus dem orthografischen, phonologischen und den zwischengeschalteten verborgenen Einheiten („hidden units“). Darüber hinaus gibt es hierbei keine Rückmeldung von den phonologischen Repräsentationen an die verborgenen Einheiten.

Die Worterkennung besteht im Wesentlichen aus der Umformung zwischen den einzelnen Repräsentationseinheiten. Das laute Lesen benötigt demnach ein orthografisches Bild, um in weiterer Folge das entsprechende phonologische Muster zu erstellen. Bei dieser Umsetzung werden zusätzlich Interaktionen der verborgenen Einheiten hinzugeschaltet, welche zwischen den orthografischen, phonologischen und semantischen Repräsentationen vermitteln. Bei den abgebildeten Modellen lässt sich eine unterschiedliche Vorgangsweise der Worterkennung feststellen. Im Gesamtmodell von Seidenberg und McClland (1989), welches in Abbildung 2 links dargestellt ist, wird von zwei möglichen Verarbeitungswegen ausgegangen: einem direkten Weg von der orthografischen zur phonologischen Einheit und einem Weg, bei dem zusätzlich semantische Repräsentationen dazwischengeschaltet sind. Bei der zweiten Möglichkeit können bei der Bedeutungsebene noch Informationen aus der Kontexteinheit hinzugezogen werden. Im Triangelmodell, das in Abbildung 2 rechts gezeigt wird, lässt sich nur eine Art der Wortverarbeitung, nämlich direkt von der orthografischen zur phonologischen Repräsentation, erkennen.

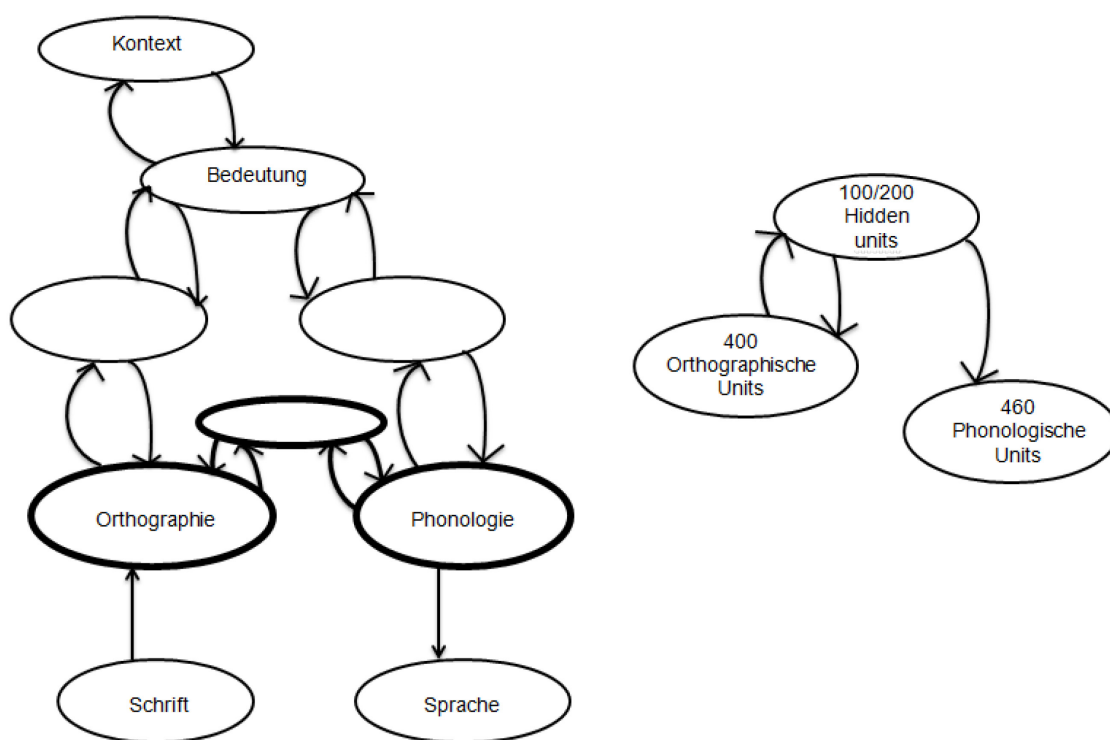


Abbildung 2: Links: Das Parallel Distributed Modell mit den phonologischen, semantischen und orthografischen Codes und dem kontextuellen Einfluss. Die Pfeile zeigen die Verbindungen zwischen den Einheiten an und die leeren Ellipsen repräsentieren die verborgenen Einheiten („hidden units“), welche zwischen den Codes vermitteln. Rechts: Der fett gedruckte Kern aus der linken Abbildung: das Triangelmodell (Grafik adaptiert aus Seidenberg & McClland, 1989).

Auf Basis dieser Annahmen wurden Computersimulationen erstellt, welche nicht nur auf regelmäßige Wörter, sondern auch auf unregelmäßige und Pseudowörter anwendbar sind. Das Dual-Route Modell von Coltheart et al. (2001) nimmt beispielsweise an, dass diese unterschiedliche Vorgangsweisen für die Aussprache benötigen. Unregelmäßige Wörter müssen nach dem Zwei-Wege-Modell aus dem mentalen Lexikon abgerufen werden, da das buchstabenweise Abrufen, wie es bei Pseudowörtern angewandt wird, zu einem falschen Output führen würde.

Im Gegensatz dazu besitzen Netzwerkmodelle, wie jenes von Seidenberg und McClelland (1989), kein mentales Lexikon, aus dem ganze Wörter abgerufen werden können. Vielmehr gehen seine VertreterInnen von nur einem Mechanismus aus. Das Netzwerk lernt nämlich mit der Erfahrung, regelmäßige und unregelmäßige Wörter, Pseudowörter und Buchstabenfolgen zu bearbeiten, indem es sich die Schreibweisen-Laut-Korrespondenzen der jeweiligen Wörter aneignet.

Das Modell berücksichtigt, wie auch das DRC Modell, nur einsilbige Wörter. Da das Netzwerkmodell mit statistischer Kovariation arbeitet, kann davon ausgegangen werden, dass Verknüpfungen zwischen Buchstaben, die häufig gemeinsam auftreten, gespeichert werden. Dadurch kann in diesem Modell eine Berücksichtigung sublexikaler Cluster beim Lesen angenommen werden.

Konnektionistisches Zwei-Wege-Modell - CDP++ (Perry, Ziegler & Zorzi, 2010)

Die bisher erwähnten Lesemodelle haben sich nur mit einsilbigen Wörtern beschäftigt. Das CDP++ Modell von Perry et al. (2010) stellt nun eine Erweiterung auf zweisilbige Wörter dar. Es kann darüber hinaus als eine Kombination der Grundannahmen der Netzwerkmodelle und der Zwei-Wege-Theorien angesehen werden.

Der Vorgang des lauten Lesens beginnt im CDP++ Modell, wie bei den bisherigen Theorien, mit einem Schriftbild, welches als Input dient. Darauf folgend werden in einem ersten Schritt dessen Eigenschaften in einer visuellen Wortanalyse bearbeitet und in weiterer Folge die Verbindungen der Buchstaben. Ähnlich wie beim DRC Modell von Coltheart et al. (2001) teilt sich das Modell in zwei Routen: dem lexikalischen und dem nicht-lexikalischen Weg. Für eine bildliche Darstellung des beschriebenen Leseablaufs siehe auch Abbildung 3.

Der sublexikale Weg erstellt die Aussprache für eine Buchstabenfolge, lässt jedoch dessen sublexikalen Status außer Acht. Des Weiteren ist er ausschlaggebend für die Entschlüsselung neuer Stimuli, wie zum Beispiel Pseudowörter. Vorerst jedoch werden die Grapheme der

Buchstabenfolge durch einen graphemischen „Parser“ analysiert und in weiterer Folge im graphemischen „Puffer“ einem entsprechenden Platz zugewiesen. Der graphemische „Puffer“ lässt sich in drei Komponenten unterteilen, dem Onset-Graphem (Beginn einer Silbe), dem Vokalgraphem (Silbenmitte) und dem Coda-Graphem (Anhang einer Silbe). Da das CDP++ Modell auch zweisilbige Wörter erlesen kann, ist dieser Aufbau in doppelter Ausführung vorhanden. Der Kern der sublexikalen Route stellt das sogenannte TLA Netzwerk (two-layer network of phonological assembly) dar – einem zweilagigen Netzwerk der Phonologiegewinnung. Dieses folgt der Aktivierung des graphemischen „Puffers“ und ordnet die Grapheme den jeweiligen Phonemen zu. Das Netzwerk erlernt während der Auseinandersetzung mit den Wörtern, die bewährtesten Verknüpfungen zwischen Orthographie und Phonologie zu erstellen.

Dieser sublexikale Weg lässt sich im Wesentlichen mit der nicht-lexikalen Route des DRC Modells von Coltheart et al. (2001), dem phonologischen Rekodieren, vergleichen. Im CDP++ Modell werden bereits erprobte Buchstaben-Laut-Zuordnungen, ähnlich dem Netzwerkmodell von Seidenberg und McClelland (1989), erlernt, wobei beim DRC Modell feste Regeln bei der Rekodierung zur Anwendung kommen.

Beim lexikalen Weg, welcher ein phonologisches und orthographisches Lexikon besitzt, werden die Repräsentationen ganzer Wörter, wie beim DRC Modell, aus dem Speicher abgerufen.

Beide Routen vereinen sich im phonologischen Output, welcher sich aus dem Phonem und dem Akzent-Output-Knoten zusammensetzt. Durch die Einbeziehung des Akzents kann das Modell feststellen, ob die Betonung auf der ersten oder der zweiten Silbe liegt. Es kann darüber hinaus nicht nur die bereits erwähnte Buchstaben-Laut-Zuordnung erlernen, sondern auch die Graphem-Akzent-Verbindung.

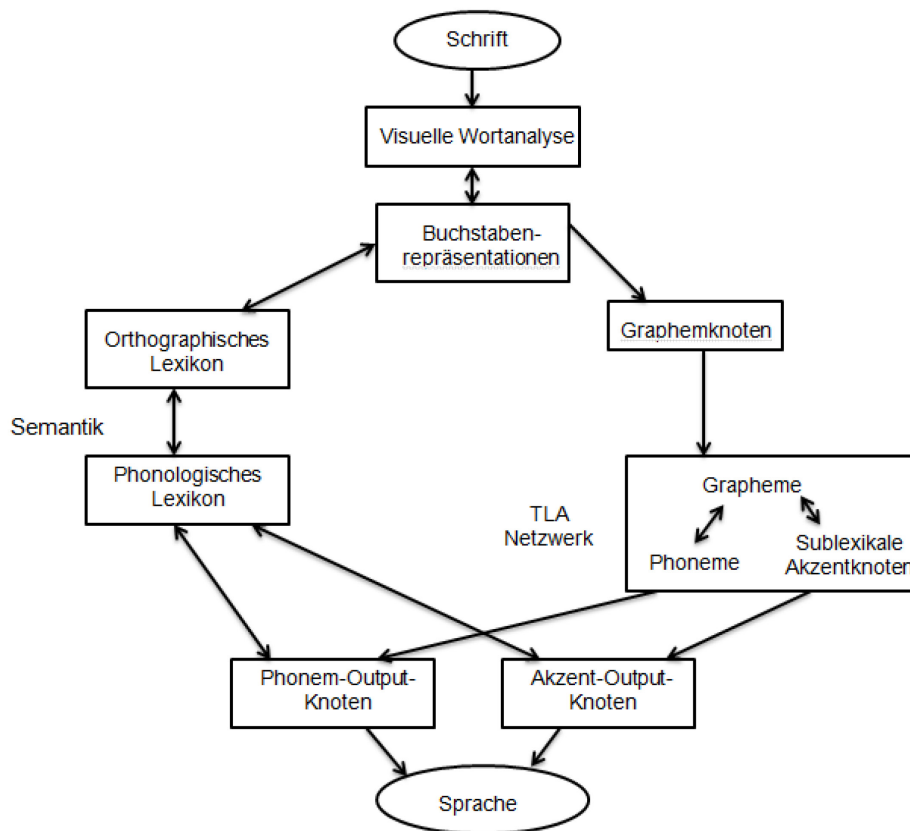


Abbildung 3: Das Konnektionistische Zwei-Wege-Modell - CDP++ (Grafik adaptiert aus Perry, Ziegler & Zorzi, 2010).

Nach Perry et al. (2010) ist die phonologische und die orthografische Einheit direkt miteinander verbunden und nicht wie beim Modell von Seidenberg und McClelland (1989), bei dem sogenannte „hidden units“ dazwischengeschaltet sind. Durch den graphemischen „Parsers“ im CDP++ Modell scheinen sublexikale Cluster ein fester Bestandteil zu sein. Die Grundannahmen von Perry et al. (2010) sprechen demnach für die Berücksichtigung sublexikaler Einheiten beim Lesen.

2.3 Bisherige Studien zur Verwendung sublexikaler Cluster

Die Verarbeitung sublexikaler Cluster beim Lesen stellt ein nicht allzu breit untersuchtes Forschungsgebiet dar. Zwar wurden einige Studien zur Verwendung von Onset und Rime als Einheit durchgeführt, die Verarbeitung von Konsonantencluster wurde jedoch meist außer Acht gelassen.

Im Folgenden sind einige Studien angeführt, die sich mit der Verarbeitung von Konsonanten- und Onset/Rime Cluster beschäftigen. Zudem werden auch Studien erwähnt, welche von einer

silbenweisen Zerlegung und Speicherung der Wörter ausgehen und diesen Sachverhalt erforschen.

Bruck, Caravolas und Treiman (1995) beschäftigten sich in zwei empirischen Studien im englischsprachigen Raum mit der Frage, ob Pseudowörter einer silbenweisen Verarbeitung unterliegen. Dabei hörten die Versuchspersonen Paare von Pseudowörtern, bei denen entweder die gesamte erste Silbe, nur ein Teil der ersten Silbe, oder keines der Phoneme identisch waren. Die ProbandInnen sollten dabei verneinen oder bejahen, ob die beiden Pseudowörter dieselben Anfangslaute besitzen. Ziel der Studie war es, zu klären, ob neben Silben und Phonemen noch weitere sublexikale Segmentierungsarten, wie beispielsweise die des Rimes, angewendet werden. Die Ergebnisse zeigten, dass die Reaktionszeiten der Versuchspersonen schneller waren, wenn die gesamte Silbe der Pseudowortpaare zu Beginn identisch war. Jene Effekte konnten bezüglich des Rimes am Ende des Wortes nicht beobachtet werden. Die Reaktionszeiten wiesen nicht signifikante Unterschiede der angeführten Versuchsbedingungen auf.

Diese Ergebnisse deuten daraufhin, dass Wörter beziehungsweise Pseudowörter silbenweise verarbeitet und segmentiert werden.

Eine der wenigen Studien, die auch die Verarbeitung von Konsonantenclustern mit einbezog, wurde von Frederiksen und Kroll (1976) durchgeführt. Sie verwendeten in ihren Benennungs- und Entscheidungsaufgaben ein- und zweisilbige Wörter und Pseudowörter, die aus unterschiedlichen Häufigkeiten und Buchstabenanzahlen bestanden.

Es wurde angenommen, dass zwei oder drei Konsonanten zu Beginn oder am Ende eines Wortes komplexere Verarbeitungsprozesse benötigen, als ein einzelner Konsonant.

In vertiefenden Analysen der orthografischen Strukturen konnte jedoch herausgefunden werden, dass, sowohl bei Wörtern als auch bei Pseudowörtern, ein größeres Konsonantencluster in der Anfangsposition (KKKVKK) mehr Zeit für den Beginn der Aussprache benötigt, als am Ende eines Wortes (KKVKKK). Eine mögliche Begründung dafür sehen Frederiksen und Kroll (1976) in der höheren Komplexität der Regeln der phonologischen Übereinstimmung von langen Konsonantenclustern.

Eines der ersten Experimente, welches mit einem visuellen Segmentierungsparadigma arbeitete, wurde von Treiman und Chafetz (1987) im englischen Sprachraum erstellt. Sie

gingen davon aus, dass Silben bei der Worterkennung in Onset und Rime gegliedert werden. Eine visuelle Störung zwischen diesen (TW IST) würde demnach das Lesen weniger beeinträchtigen, als eine Trennung im Onset einer Silbe beziehungsweise in dessen Rime (TWI ST).

In einem Experiment wurden Buchstabenreihen, bestehend aus zehn Großbuchstaben, vorgegeben. Im positiven Durchgang konnte aus Buchstabengruppen ein Wort gebildet werden, wobei dies im negativen Durchgang nicht möglich war. Als visuelle Störung wurde ein Leerzeichen entweder nach einem Konsonantencluster zwischen Onset und Rime (zum Beispiel: FL ANK) oder im Rime (zum Beispiel: FLA NK) eingebaut.

Die Ergebnisse zeigten, dass Wörter schneller aus der Buchstabenkette erkannt werden konnten, wenn eine Trennung zwischen Onset und Rime vorgenommen wurde, als bei einer Trennung im Rime.

In einem weiteren Experiment wurde ein visuelles Störsymbol (/) entweder nach einem Konsonantencluster zwischen Onset und Rime (CR//ISP) oder im Rime (CRI//SP) in Wörter und Pseudowörter eingebaut. Die Versuchspersonen sollten dabei entscheiden, ob es sich um ein Wort handelte oder nicht und den entsprechenden „ja“ oder „nein“ Knopf betätigen. Es konnte herausgefunden werden, dass die Entscheidungsaufgaben schneller getroffen wurden, wenn die Trennung zwischen Onset und Rime eingefügt wurde.

Diese Ergebnisse lassen darauf schließen, dass orthografische Einheiten dieser Art Bestandteil der Leseverarbeitung sind.

Eine weitere englischsprachige Untersuchungsreihe, die eine visuelle Segmentierung in Wörter einbaute, erstellten Levitt, Healy und Fendrich (1991). Dabei verfassten sie zwei Arten von einsilbigen Wörtern und Pseudowörtern, die ein Konsonantencluster zu Beginn und am Ende aufwiesen (KKVKK): eine, um die Effekte der Störung (*) im Konsonantencluster zwischen Onsets und Rime (zum Beispiel: CR*AFT) zu untersuchen und eine, um dies im Cluster zwischen Kern (dem Vokal in der Mitte des Wortes) und Silbenauslaut (zum Beispiel: BLU*NT) zu erforschen. Mit dieser Studie wollten sie das Vorhandensein einer silben-internen Struktur und deren Auswirkungen auf die Aussprache erforschen. Es konnte festgestellt werden, dass sich ein Rime in weitere Einheiten gliedern lässt und dass hierbei ein Konsonant, der einem Vokal folgt, Auswirkungen auf die Aussprache hat. Weiters konnte jedoch kein Effekt bei einer Störung zwischen Onset und Rime beobachtet werden. Entgegen

diesen Erwartungen kann nach diesen Ergebnissen nicht von einer Teilung einer Silbe in Onset und Rime ausgegangen werden.

Auch Bowey (1996) arbeitete im englischen Sprachraum in einer Versuchsreihe mit einer visuellen Segmentierung der Wörter. Als Störung diente hierbei die Präsentation der Stimuli in Groß- und Kleinbuchstaben (zum Beispiel: chuRCH), welche laut vorgelesen werden sollten.

In Experiment 1 und 2 wurde in der Experimentalbedingung der Onset gestört („SHred“) und in der Kontrollbedingung blieb dieser vollständig („SHawl“). Die Annahme, dass die Wörter mit einer visuellen Störung im Onset langsamer gelesen werden, als jene in der Kontrollbedingung, konnte hierbei bestätigt werden.

Auch in Experiment 3 wurden die Auswirkungen der visuellen Störung des Onsets beim Lesen untersucht. Dafür wurde in der Experimentalbedingung der erste Buchstabe klein und die verbleibenden groß gezeigt („tANK“). In der Kontrollbedingung wiederum wurden die Wörter in Großbuchstaben präsentiert („FLAG“). Entgegen den Erwartungen konnte hier kein Störeffekt nachgewiesen werden.

In Experiment 4 wurde die visuelle Segmentierung in der Versuchsbedingung im Rime eingebaut („buMP“) und in der Kontrollbedingung blieb dieser vollständig („trOT“). Die Annahme, dass die Wörter mit einem gestörten Rime langsamer gelesen werden, als die ungestörten, bestätigte sich.

Experiment 5 befasste sich ebenfalls mit der visuellen Segmentierung des Rimes. Dabei wurden die Wörter in der Experimentalbedingung zu Beginn groß und am Ende klein gezeigt („Dent“). In der Kontrollbedingung wurde die Buchstabenfolge in Großbuchstaben präsentiert („BRED“). Hier konnte die Annahme, dass Wörter in der Experimentalbedingung langsamer gelesen werden, als jene in der Kontrollbedingung, ebenfalls bestätigt werden.

Obwohl eine Versuchsanordnung zu nicht signifikanten Ergebnissen führte, kann trotz allem von einer Verarbeitung des Onset und Rime als Einheit beim Lesen ausgegangen werden.

Im niederländischen Sprachraum befassten sich Martensen, Maris und Dijkstra (2003) unter anderem mit Leseaufgaben von Wörtern, die durch ein Symbol (/) gestört wurden. Der Fokus dieser Studie lag in der Auswirkung der visuellen Trennung von zwei Buchstaben, die gemeinsam als Einheit bei der Aussprache verarbeitet werden. Dabei wurde einerseits ein vokalischer Digraph (zum Beispiel: BO//EK) und andererseits der Buchstabe „C“ des Onsets

und der darauffolgende Vokal, der für dessen Betonung wesentlich ist (beispielsweise bei C//ORTEX), getrennt.

Die Ergebnisse zeigten, dass die Störung von zwei Vokalen das Leseverhalten in allen Versuchsbedingungen beeinträchtigte und diese scheinbar als Cluster verarbeitet werden. Die Trennung des Buchstaben „C“ im Onset von dem darauffolgenden Vokal hatte keine größeren Auswirkungen auf das Lesen, als eine Störung nach dem zweiten Buchstaben (wie zum Beispiel: CO//RTEX). Martensen et al. (2003) folgerten, dass Buchstaben, die ein Phonem bilden, nicht als Einheit beim Lesen verarbeitet werden.

Marinus und de Jong (2008) wollten in ihrer Studie herausfinden, ob Buchstabencluster beim Lesen als Einheit verarbeitet werden. In diesem Sinn führten sie einerseits ein Trennungssymbol im Konsonantencluster des Onsets (s#top) und andererseits ein Symbol zwischen Onset und Rime (st#op) ein. Um Positionseffekte zu vermeiden, erstellten sie eine visuelle Störung im Rime. Darüber hinaus trennten sie ebenfalls den vokalischen Digraph im Wort. Daraus ergaben sich insgesamt vier Versuchsbedingungen: (a) Konsonantencluster zu Beginn gestört (C#CVVC), (b) Trennungssymbol im Vokal-Digraph (CCV#VC), (c) Rime gestört (CCVV#C) und (d) Wort ohne eingebautes Störsignal (#CCVVC). Sie nahmen an, dass, wenn ein Cluster beim Lesen als Einheit verarbeitet wird, die Einführung des Störsignals im Cluster des Onsets zu einer stärkeren Verlangsamung der Lesegeschwindigkeit führt, als wenn das Trennungssymbol zwischen Onset und Rime steht und das Cluster im Anlaut vollständig bleibt.

Marinus und de Jong (2008) führten diese Studie mit durchschnittlich begabten LeserInnen und mit legasthenen Kindern der vierten Schulstufe durch. Die Ergebnisse zeigten, dass die Einführung einer visuellen Trennung im Onset oder im Rime einen ähnlichen Effekt sowohl bei begabten als auch bei weniger begabten LeserInnen, hat als ein Störsignal zwischen Onset und Rime. Daraus lässt sich folgern, dass die Kinder Onset und Rime nicht als Einheit verarbeiten. Weiters zeigten sich beide Gruppen bei der Trennung der Vokal-Digraphie am meisten beeinträchtigt. Dies spricht dafür, dass die Kinder wiederum diese Clusterart als Einheit beim Lesen verwenden.

Um auszuschließen, dass sich diese sublexikalen Cluster noch nicht vollständig bei den SchülerInnen etabliert haben, wurden in einer darauffolgenden Studie Erwachsene als Versuchspersonen herangezogen. Letztendes wiesen diese, entgegen den Erwartungen, keine Verringerung der Lesezeit bei der visuellen Trennung des Konsonanten-Onsets, verglichen

mit der Störung zwischen Onset und Rime, auf. Im Gegensatz zu den Kindern zeigten diese Versuchspersonen jedoch, dass das Trennungssymbol im Konsonantencluster des Onsets einen geringeren negativen Effekt auf die Lesezeit hatte, als die visuelle Störung außerhalb der Einheit. Diese Versuchsreihe lässt wiederum den Schluss zu, dass sublexikale Cluster beim Lesen nicht als Einheit verarbeitet werden.

Ein sogenanntes Priming Experiment, bei dem implizite Gedächtnisinhalte durch vorhergehende Reize („primes“) aktiviert werden, erstellte unter anderem Bowey (1990) in Form einer mehrteiligen Versuchsreihe. Sie primte sowohl den Onset, als auch den Rime eines Wortes, indem die letzten, oder die ersten Buchstaben als Prime fungieren sollten. In der Kontrollbedingung wurden den VersuchsteilnehmerInnen anstelle des Priming-Reizes Sternsymbole gezeigt. Es konnte in allen Bedingungen ein signifikanter Haupteffekt der Aktivierung festgestellt werden: Worte, bei denen zuvor ein Priming-Reiz präsentiert wurde, wurden schneller benannt als jene in der Kontrollbedingung.

Die Ergebnisse lassen laut Bowey (1990) darauf schließen, dass eine Buchstabenfolge vorerst in Onset und Rime zergliedert und erst im weiteren Verlauf synthetisiert wird. Die Worterkennung geht bei dieser Sichtweise nicht von einem lexikalen Mechanismus, wie es beispielsweise beim DRC Modell von Coltheart et al. (2001) der Fall ist, aus.

Ein weiteres Priming-Experiment wurde von Grainger & Ferrand (1996) durchgeführt. Die Autoren konnten nachweisen, dass die Wortbenennung vereinfacht wurde, wenn der Prime und der Ziel-Stimulus den gleichen Onset hatten.

Mousikou, Coltheart, Saunders und Yen (2010) beschäftigten sich in ihrer Studie mit der Verarbeitung von Onset und Rime eines Wortes beim lauten Lesen. Es sollte dabei herausgefunden werden, ob diese Cluster während des Leseprozesses als Einheit oder als einzelne Buchstaben repräsentiert werden.

In diesem Zusammenhang kamen verschiedene Priming-Experimente zur Anwendung. In jedem davon sollte getestet werden, ob die Ziel-Onsets schneller benannt werden, wenn der zuvor präsentierte Prime denselben Anfangsbuchstaben wie das Zielwort hat („disc-drum“; verwandte Bedingung), als wenn Prime und Wort nicht denselben Anfangsbuchstaben besitzen („melt-drum“; nicht-verwandte Bedingung). Mousikou et al. (2010) erstellten zur

Untersuchung dieser Annahmen einsilbige aussprechbare Pseudowörter, die aus vier Buchstaben bestanden.

Die Studien zeigten, dass die Wortbenennung in der verwandten Bedingung schneller erfolgte, als dies in der unverwandten der Fall war.

Beide Pseudowortarten teilen sich in dieser Versuchsreihe nur den ersten Buchstaben, nicht den gesamten Anlaut (Onset). Da die Wörter in der verwandten Bedingung rascher benannt werden konnten als in der unverwandten Bedingung, lässt sich daraus folgern, dass das menschliche Lesesystem den Onset nicht als Einheit verarbeitet, sondern ihn in einzelne Buchstaben zerlegt. Hätte man keine Unterschiede zwischen den Gruppen beobachten können, ließe sich auf eine Bearbeitung des Onsets als Einheit rückschließen. Da die gefundenen Ergebnisse nur auf Pseudowörter zutreffend sind, wurde unter denselben Versuchsanordnungen eine Testreihe mit einsilbigen Wörtern erstellt. Auch hier traten die gleichen Ergebnisse wie zuvor ein: Wörter in der verwandten Bedingung konnten schneller benannt werden, als in der unverwandten Bedingung. Dies spricht auch hier gegen eine Verarbeitung der Onsets als Cluster.

Eine Versuchsreihe, die sich mit der Frage befasste, ob Bilder oder Wörter schneller benannt werden, wurde von Hennessey und Kirsner (1999) durchgeführt. Zusätzlich dazu wurde untersucht, ob dabei auch ein Effekt der Worthäufigkeit nachgewiesen werden kann. Es konnte gezeigt werden, dass Bilder insgesamt schneller benannt werden als Wörter, es jedoch auch einen Zusammenhang mit der Worthäufigkeit gibt. Dieser Effekt zeigte sich jedoch nur bei Wörtern mit geringer Häufigkeit. Sublexikale Prozesse sollen hierbei vermehrt involviert sein, als bei hochfrequenten Wörtern. Jene, die eine hohe Häufigkeit besitzen, werden demzufolge vorwiegend auf Basis des lexikalen Übertrags benannt.

Im Sinne des Dual-Route Modells überträgt der sublexikalische Weg Phoneme stückweise zur Produktion. Nach dem Autor/der Autorin wird der Onset eines Wortes ausgesprochen, obwohl der verbleibende Rest noch nicht vollständig entschlüsselt wurde.

Die berichteten Forschungsergebnisse stellen einen interessanten Ausgangspunkt für weitere Untersuchungen dar. Da die vorgestellten Studien entweder im englischen oder niederländischen Sprachraum durchgeführt wurden, sind diese aus diesem Grund nicht vorbehaltlos in das Deutsche übertragbar.

Die beschriebenen Ergebnisse lassen zudem keinen eindeutigen Rückschluss zu, ob Onset beziehungsweise Rime und Konsonantencluster als Einheit beim Lesen verarbeitet werden. Vor allem die Verwendung von Konsonantencluster stellt ein wenig erforschtes Gebiet im Bereich der Leseforschung dar.

Marinus und de Jong (2008) arbeiteten in ihrer Studie mit einer symbolischen visuellen Trennung im Konsonantencluster des Onsets. Sie konnten dabei herausfinden, dass dies jedoch keinen Einfluss auf den Lesefluss darstellt. Konsonantencluster zu Beginn eines Wortes scheinen demnach nicht als Einheit verarbeitet zu werden. Da sich in der Studie der vorliegenden Arbeit das Konsonantencluster im Rime des Wortes befindet, lässt sich keine eindeutige Verbindung zur Untersuchung von Marinus und de Jong (2008) herstellen.

Die Ergebnisse der Priming-Experimente von Bowey (1990) sprechen für eine Verwendung von Onset und Rime als Cluster, wobei die Priming-Studien von Mousikou et al. (2010) wiederum dagegen sprechen und demnach eine serielle Verarbeitung angenommen wird. Jene Untersuchungen, die mit einer visuellen Segmentierung in Form von Symbolen (* oder //) gearbeitet haben, lassen dagegen einheitlichere Rückschlüsse zu: Levitt et al. (1991), Martensen et al. (2003) und Marinus und de Jong (2008) gehen aufgrund ihrer Ergebnisse davon aus, dass Onset und Rime nicht als Einheit verarbeitet werden. Bowey (1996) trennte Wörter, indem sie dessen Wortbild in Groß- und Kleinschreibung präsentierte (buMP). Diese Art der visuellen Trennung widersprach den zuvor erwähnten Ergebnissen und zeigte, dass sowohl Onset, als auch Rime als Einheit fungieren.

In der Studie der vorliegenden Arbeit wird, ähnlich wie bei Bowey (1996), eine visuelle Segmentierung in Form von Groß- und Kleinschreibung in das Wortbild eingeführt (zum Beispiel bunT). Die Trennung soll dabei helfen herauszufinden, ob sowohl Onset und Rime, als auch Konsonantencluster als Einheiten beim lauten Lesen verarbeitet werden. Dabei wird angenommen, dass eine Störung im Cluster den Leseprozess erschwert und in weiterer Folge verlangsamt. Marinus und de Jong (2008) führten in ihrer Untersuchung Überlegungen der Art dieser Störung zwar an, ließen sie jedoch bei der weiteren Durchführung unberücksichtigt. Baron und Strawson (1976) zeigten in einer Studie, dass eine alternierende Schrift im Wort („baUCH“) dessen Erkennung und Verarbeitung beeinflusst. Es wird dem Leser/der Leserin somit erschwert, Buchstabeneinheiten zu erfassen und es muss aus diesem Grund eine

buchstabenweise Rekodierung angewendet werden. Nicht nur die Befunde von Baron und Strawson (1976), sondern auch jene von Bowey (1996) lassen den Rückschluss zu, dass eine visuelle Segmentierung in Form von wechselnder Groß- und Kleinschreibung als geeignete Störung dient, um brauchbare Ergebnisse in Bezug auf die Verwendung sublexikaler Cluster beim Lesen zu erlangen.

Diese Untersuchung beschäftigt sich, wie bereits in den einleitenden Worten erwähnt wurde, einerseits mit der Verarbeitung von Onset und Rime und andererseits mit jener von Konsonantenclustern als Einheit, wobei die dargebotenen Wörter und Pseudowörter laut vorgelesen werden sollen. Dabei ist in einer Bedingung das betroffene sublexikale Cluster gestört und in der zweiten Bedingung ist dieses ungestört.

Als Kontrollbedingung werden auch Wörter und Pseudowörter vorgegeben, bei denen der Rime ungestört bleibt und die Segmentierung zwischen Onset und Rime erfolgt. Auch die Konsonantencluster werden sowohl gestört, als auch ungestört in der Kontrollbedingung dargeboten.

Wenn jene sublexikalen Cluster beim lauten Lesen als Einheiten verarbeitet werden, müssten Wörter und Pseudowörter, die eine Trennung darin aufweisen, langsamer gelesen werden, als jene, dessen Einheiten vollständig bleiben. Die Fragestellung der vorliegenden Arbeit behauptet demnach, dass Wörter und Pseudowörter, welche eine visuelle Segmentierung im Rime oder im Konsonantencluster haben, langsamer gelesen werden, als diejenigen in den Kontrollbedingungen, bei denen die Einheiten ungestört bleiben.

Wie bereits in den Entwicklungsmodellen des Lesens beschrieben wurde, werden mit steigender Lesefertigkeit zunehmend Buchstabencluster gebildet und auch als solche gespeichert, um den Leseprozess zu vereinfachen und zu beschleunigen. Diese gesteigerte Lesefertigkeit kann erst gegen Ende der zweiten Klasse erwartet werden. Grund dafür ist der österreichische Lehrplan für Volksschulen, worin festgelegt ist, dass SchülerInnen erst zu diesem Zeitpunkt flüssig lesen können sollen. Die 2. und die 4. Schulstufe stellen demnach eine geeignete Grundlage zur Stichprobenziehung dar. Da die Untersuchung zu Schulbeginn durchgeführt wurde, kann gut ein Vergleich zwischen LeseanfängerInnen (2. Schulstufe) und fortgeschrittenen LeserInnen (4. Schulstufe) gezogen werden. Ein Trennungseffekt könnte sich aufgrund der zuvor erwähnten Argumente erst bei LeserInnen der 4. Schulstufe zeigen.

Um darüber hinaus auch einen Zusammenhang zwischen der Lesekompetenz und der Verwendung sublexikaler Cluster analysieren zu können, wurde in der durchgeführten Studie pro Schulstufe zwischen guten und schlechten LeserInnen differenziert.

3. Untersuchung

Die Untersuchung der vorliegenden Arbeit gliederte sich in eine Vor- und eine Haupterhebung. In der ersten Phase wurden die Leseleistungen der SchülerInnen mittels Salzburger Lesescreening (Mayringer & Wimmer, 2003) erfasst und somit die 22 besten und schlechtesten LeserInnen pro Schulstufe ermittelt. Diese wiederum nahmen in der zweiten Phase an einem selbst konstruierten Wortlesetest teil.

3.1. TeilnehmerInnen

Es wurden insgesamt 232 Elternbriefe an SchülerInnen der zweiten und vierten Schulstufe an einer Wiener und drei burgenländischen Schulen verteilt.

Die Erlaubnis für die Durchführung wurde sowohl von Stadt- als auch vom burgenländischen Landesschulrat und von den jeweiligen Direktorinnen und LehrerInnen eingeholt. Auch die Eltern wurden durch diesen Elternbrief über den Ablauf der Erhebung informiert und gebeten, eine Einverständniserklärung zu unterschreiben. Die Teilnahme war freiwillig und darüber hinaus wurden alle Daten der SchülerInnen vertraulich behandelt und anonymisiert.

Aufgrund von nicht vorhandenen Einverständniserklärungen konnten nur 179 von den 232 SchülerInnen in die Erhebung mit einbezogen werden.

3.2. Instrumente

Vorerhebung: Salzburger Lesescreening (Mayringer & Wimmer, 2003)

Die individuelle Leseleistung jedes Schülers/jeder Schülerin wurde durch das Salzburger Lesescreening (Mayringer & Wimmer, 2003) erfasst, welches in österreichischen Schulen regelmäßig angewendet wird. Dieses Verfahren ist ein Gruppenverfahren, bei dem die Kinder

verschiedene Sätze mit steigender Länge lesen und durch Einkreisen bewerten sollen, ob diese wahr oder falsch sind. Das Häkchen stand für eine richtige und das Kreuz für eine falsche Aussage.

Einige Beispiele dafür sind:

Kirschen können sprechen.	✓	✗
Eine Woche hat sieben Tage.	✓	✗
Vögel bauen Nester.	✓	✗

Abbildung 4: Beispiele der zu bearbeitenden Sätze aus dem Salzburger Lesescreening (Mayringer & Wimmer, 2003).

Als Testwert diente die Anzahl der in drei Minuten richtig bewerteten Sätze. Davor konnten sich die Kinder in einem einminütigen Übungsdurchgang an den Verfahrensmodus gewöhnen.

Es wurden insgesamt 12 Kinder aufgrund ihrer Ergebnisse im Lesescreening von der Haupterhebung ausgeschlossen, da sie in der 2. Schulstufe weniger als zehn Sätze und in der 4. Schulstufe weniger als 25 Sätze richtig erlasen. Durch diese Resultate musste davon ausgegangen werden, dass diese SchülerInnen im Wortlesetest der Haupterhebung größere Schwierigkeiten haben würden.

Jene Kinder, die ein verwertbares Ergebnis erreichten, erlasen innerhalb der drei Minuten in den zweiten Klassen (n=80) durchschnittlich 22 Sätze (M=22,4; SD=6,77) und in den vierten Klassen (n=86) durchschnittlich 41 Sätze (M=40,71; SD=8,7) richtig.

Davon wurden die jeweils 22 schnellsten und langsamsten LeserInnen pro Schulstufe in die Haupterhebung aufgenommen. Die schnellsten LeserInnen der 2. Schulstufe lasen durchschnittlich 31 Sätze (M=30,86; SD=4,93) richtig, die langsamsten wiederum durchschnittlich 15 Sätze (M=14,91; SD=1,93). Die schnellsten LeserInnen der 4. Schulstufe erlasen durchschnittlich 52 Sätze (M=51,86; SD=7,27) richtig und die langsamsten durchschnittlich 31 Sätze (M=31,14; SD=2,27) richtig.

Haupterhebung: Wortlesetest

Der Wortlesetest bestand aus einsilbigen Wörtern und Pseudowörtern aus vier Buchstaben, welche eine visuelle Trennung in Form von alternierender Groß- und Kleinschreibung beinhalteten und von den Kindern laut vorgelesen werden sollten.

Für ihre Erstellung wurden Silben mit durchschnittlicher und hoher Häufigkeit aus der CELEX-Datenbank (Baayen, Piepenbrock & van Rijn, 1993) herausgesucht und damit mit einem Wortgenerierungsprogramm Wörter und Pseudowörter erstellt. Die Erhebung konzentrierte sich einerseits auf die Verarbeitung von Onset und Rime und andererseits auf jene von Konsonantencluster.

Um auch jene Konsonantencluster zu erforschen, die im Schulkontext öfter vorkommen, wurden, wie bereits erwähnt wurde, neben den durchschnittlichen auch hohe Silbenhäufigkeiten mit einbezogen. Die logarithmierten Häufigkeiten der generierten Wörter lagen insgesamt zwischen 1 und 2.9943.

Mit diesen Silben wurden insgesamt sechs Trennungsbedingungen aus jeweils fünf Wörtern und Pseudowörtern erstellt, welche den SchülerInnen vorgegeben wurden:

- (a) Trennung zwischen Onset und Rime (sEIL),
- (b) Trennung im Rime (seiL),
- (c) keine Trennung eines Konsonantenclusters durchschnittlicher Häufigkeit (heLD),
- (d) Trennung eines Konsonantenclusters durchschnittlicher Häufigkeit (helD),
- (e) keine Trennung eines Konsonantenclusters hoher Häufigkeit (baCH) und
- (f) Trennung eines Konsonantenclusters hoher Häufigkeit (bach).

Eine genaue Auflistung des Stimulusmaterials befindet sich im Anhang A.

Jedes Wort beziehungsweise Pseudowort wurde den Kindern einmal präsentiert, wobei bei jeder Trennungsbedingung zuerst die Wörter und anschließend die Pseudowörter vorgegeben wurden.

2.3. Durchführung

Der Wortlesetest wurde den SchülerInnen einzeln während des regulären Schulbetriebs in einem separaten und ruhigen Raum vorgegeben. Den Kindern wurde vorerst erklärt, dass sie

die folgenden Wörter rasch und richtig laut vorlesen sollten. Sie wurden weiters darauf hingewiesen, dass die Groß- und Kleinbuchstaben in den Wörtern vertauscht, beziehungsweise durcheinander gekommen sein könnten. Diese Tatsache sollten die SchülerInnen nicht beachten und sich auf das korrekte Lesen des Wortes konzentrieren. Bei den Pseudowörtern wurden die Kinder darauf aufmerksam gemacht, dass ebenfalls Buchstabenfolgen vorkommen können, die im Grunde keine richtigen Wörter sind, jedoch genauso behandelt und vorgelesen werden sollten. Vor der eigentlichen Testung wurden insgesamt vier Übungsitems vorgegeben. Jeder Durchgang dauerte inklusive der Instruktion etwa 10 Minuten.

Das Stimulusmaterial wurde auf einem 15,7 Zoll Bildschirm eines Laptops präsentiert, wobei die SchülerInnen etwa 50 cm von diesem entfernt saßen. Die Wörter und Pseudowörter wurden in einer PowerPoint Präsentation in der Schriftart Leelawadee in Schriftgröße 75, schwarz geschrieben, gezeigt. Vor und nach jedem Wort beziehungsweise Pseudowort folgte ein Fixierungspunkt ($\oplus$), um den Kindern die Orientierung am Bildschirm zu erleichtern. Die SchülerInnen wurden während der Testdurchführung mit einem Audioprogramm über den Laptop aufgenommen. Nach jeder Folie ertönte ein Klick, um die Tonspur eines neu gezeigten Wortes beziehungsweise Pseudowortes im Audioprogramm zu markieren. Die Testleiterin schaltete selbst per Mausklick die Folien weiter und die Leseleistungen der SchülerInnen wurden durch das Mikrofon des Laptops mit dem Aufnahmeprogramm für die spätere Auswertung aufgezeichnet.

Kodierung

Es wurden sowohl die Latenzzeiten als auch die Gesamtnennzeiten jedes Wortes und Pseudowortes erhoben und in die SPSS Datenmaske eingegeben.

Darüber hinaus wurden die Wörter und Pseudowörter dahingehend kodiert, ob sie spontan richtig (Kodierung 0), zuerst falsch, jedoch nach Selbstkorrektur richtig (Kodierung 1), vorerst falsch und auch nach der Selbstkorrektur falsch (Kodierung 2) oder sofort falsch und ohne Korrektur (Kodierung 3) erlesen wurden.

Auf falsche Betonung oder Akzentsetzung wurde nicht geachtet und jene Items wurden als korrekt kodiert.

Nach Durchsicht der Daten wurde entschieden, nur jene Wörter und Pseudowörter in die Auswertung mit einzubeziehen, die richtig erlesen wurden. Darüber hinaus wurde lediglich mit den jeweiligen Latenzzeiten weitergerechnet.

4. Ergebnisse

Zuerst wurde der Durchschnittswert der Latenzzeit pro Trennungsbedingung berechnet. Dabei wurden, wie bereits zuvor erwähnt, nur jene Wörter und Pseudowörter mit einbezogen, die richtig erlesen wurden.

Die durchschnittlichen Latenzzeiten wurden einer Varianzanalyse mit Messwiederholung unterzogen. Dabei dienten die Lesekompetenz (Gruppe: gute/schlechte LeserInnen) und die Schulstufe als Zwischensubjektfaktoren und die Trennungsbedingungen als Innersubjektfaktoren.

Für eine detaillierte Analyse wurden gute und schlechte LeserInnen pro Schulstufe getrennt betrachtet. Kontraste wurden für die Prüfung der Signifikanz der Innersubjektfaktoren angefordert.

Für eine bessere Übersicht wurden die Konsonantencluster und die Onset/Rime Cluster getrennt analysiert.

Wörter und Pseudowörter wurden in jedem Fall getrennt voneinander dargestellt. Die folgenden Ergebnisse werden mit einem Signifikanzniveau von $\alpha=0,01$ berichtet.

Die durchschnittlichen Latenzzeiten der SchülerInnen von Wörtern und Pseudowörtern und deren Standardabweichungen pro Schulstufe, aufgeteilt in gute/schlechte LeserInnen sind in Tabelle 1 und 2 dargestellt.

Wörter

	2.Schulstufe		4.Schulstufe	
	gute LeserInnen	schlechte LeserInnen	gute LeserInnen	schlechte LeserInnen
Trennungsbedingung				
onset/rime	1.163 (.374)	1.548 (1.007)	.677 (.193)	1.149 (.359)
onsetri/me	1.133 (.470)	1.620 (1.119)	.828 (.428)	1.203 (.439)
Konsonantencluster ungestört (mittlere Häufigkeit)	.952 (.265)	1.377 (.673)	.648 (.153)	1.121 (.519)
Konsonantencluster gestört (mittlere Häufigkeit)	.999 (.253)	1.492 (.740)	.635 (.184)	.977 (.454)
Konsonantencluster ungestört (hohe Häufigkeit)	.982 (.286)	1.330 (.625)	.662 (.141)	.985 (.274)
Konsonantencluster gestört (hohe Häufigkeit)	1.007 (.334)	1.476 (.766)	.711 (.147)	1.070 (.405)

Tabelle 1: Durchschnittliche Latenzzeiten der Wörter (in Sekunden) pro Schulstufe, aufgeteilt in gute/schlechte LeserInnen. Standardabweichungen in Klammern.

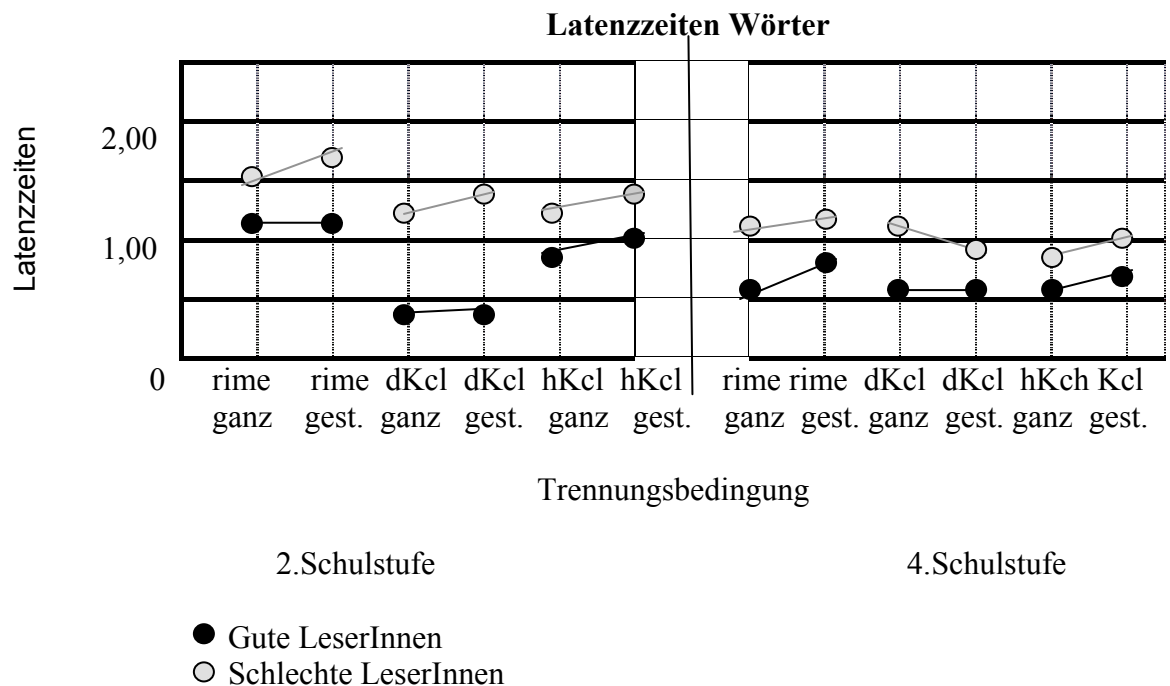


Abbildung 5: Durchschnittliche Latenzzeiten der Wörter pro Schulstufe in Sekunden, aufgeteilt in gute/schlechte LeserInnen.

Es ließ sich ein Haupteffekt der Gruppe beobachten, $F(3, 84) = 10.69$, $p = .000$, $r = .28$. Jene SchülerInnen, die in der Vorerhebung langsamer lasen, erzielten demnach auch in der Haupterhebung höhere Latenzzeiten.

Es gab weiters auch einen Haupteffekt der Schulstufe. SchülerInnen der 4. Klassen lasen durchschnittlich schneller, als jene der 2. Klassen.

Paarweise Vergleiche der Gruppen auf Ebene der Schulstufen zeigten, dass die Lesekompetenzen der guten LeserInnen der 2. Schulstufe und den schlechten LeserInnen der 4. Schulstufe annähernd gleich sind, $p = .745$.

Bei den angeforderten Interaktionseffekten war die Sphärizität verletzt, $p = .000$, und deshalb mussten die Freiheitsgrade mittels Huynh-Feldt Schätzer korrigiert werden, $\epsilon = .839$.

Es zeigte sich ein Interaktionseffekt zwischen Clusterart, Störung und Schulstufe, $F(2, 144.285) = 4.04$, $p = .026$, $r = .045$. Kontraste zeigten, dass Wörter mit gestörtem Konsonantencluster durchschnittlicher Häufigkeit in der 4. Schulstufe signifikant schneller gelesen wurden, als dies in der 2. Schulstufe der Fall war, $F(1, 86) = 5.15$, $p = .026$, $r = .057$.

Analysen der einzelnen Gruppen

2. Schulstufe

Gute LeserInnen

In dieser Gruppe konnte kein Haupteffekt der Störung beobachtet werden.

Schlechte LeserInnen

Hierbei konnte ein übergreifender Haupteffekt der Störung beobachtet werden, $F(1, 21) = 4.67$, $p = .042$, $r = .182$. Wörter mit gestörtem Cluster wurden in jeder Clusterart langsamer erlesen als jene mit ungestörtem. Die angeforderten Kontraste bestätigten diese Annahme.

Analyse beider Gruppen

Konsonantencluster

Bei einer Aufteilung nach Gruppen zeigte sich, wie zuvor angeführt, nur bei den schlechten LeserInnen ein Haupteffekt der Störung. Wurden beide Gruppen (gute/schlechte LeserInnen) zusammengefasst und diese hinsichtlich eines Störeffekts analysiert, konnte bei beiden

Häufigkeiten der Konsonantencluster ein signifikanter Haupteffekt beobachtet werden. Dies bedeutet, dass, unabhängig von der Lesekompetenz, ungestörte Konsonantencluster durchschnittlicher Häufigkeit, $F(1, 41) = 6.19$, $p = .017$, $r = .129$, als auch jene mit hoher Häufigkeit, $F(1, 42) = 5.41$, $p = .025$, $r = .114$, schneller gelesen wurden als gestörte Konsonantencluster.

4. Schulstufe

Bei der Analyse der einzelnen Gruppen, konnten weder bei den guten LeserInnen, noch bei den schlechten LeserInnen signifikante Haupteffekte der Störung beobachtet werden.

Analyse beider Gruppen

Bei einer Aufteilung nach Gruppen zeigten sich, wie zuvor erwähnt, keine Haupteffekte der Störung. Wurden beide Gruppen (gute/schlechte LeserInnen) zusammengefasst und diese hinsichtlich eines Störeffekts analysiert, konnte sowohl bei den Onset/Rime als auch bei den Konsonantenclustern signifikante Haupt- und Interaktionseffekt beobachtet werden.

Onset/Rime Cluster

Hier konnte ein Haupteffekt der Störung beobachtet werden, $F(1, 42) = 4.13$, $p = .049$, $r = .089$. Kontraste verdeutlichten, dass Wörter mit einem ungestörten Onset/Rime Cluster, unabhängig von der Lesekompetenz, schneller gelesen wurden, als Wörter mit einem gestörten Cluster.

Konsonantencluster

Bei beiden Häufigkeiten der Konsonantencluster zeigte sich ein Haupteffekt der Störung, jedoch in gegensätzlicher Richtung.

Es zeigte sich, dass Wörter, die ein gestörtes Konsonantencluster durchschnittlicher Häufigkeit besaßen, unabhängig von der Lesekompetenz, schneller gelesen wurden als Wörter mit einem ungestörten Cluster $F(1, 42) = 16.149$, $p = .000$, $r = .278$.

Auch bei Konsonantencluster hoher Häufigkeit konnte ein Haupteffekt der Störung beobachtet werden, $F(1, 42) = 5.41$, $p = .025$, $r = .114$. Kontraste zur Prüfung der Signifikanz zeigten, dass Wörter mit ungestörtem Konsonantencluster, unabhängig von der Lesekompetenz, schneller gelesen wurden als Wörter mit gestörtem Cluster.

Zudem gab es einen Interaktionseffekt zwischen dem Cluster und der Gruppe,

$F(1, 42) = 11.03, p = .002, r = .208$. Die angeforderten Kontraste zeigten, dass schlechte LeserInnen Wörter mit ungestörtem Cluster signifikant langsamer erlasen als gute LeserInnen.

Pseudowörter

	2.Schulstufe		4.Schulstufe	
	gute LeserInnen	schlechte LeserInnen	gute LeserInnen	schlechte LeserInnen
Trennungsbedingung				
onset/rime	1.516 (.692)	1.713 (.840)	1.068 (.550)	1.493 (.647)
onsetri/me	1.288 (.539)	1.636 (.888)	1.007 (.440)	1.400 (.506)
Konsonantencluster ungestört (mittlere Häufigkeit)	1.010 (.290)	1.398 (.725)	.752 (.248)	1.188 (.517)
Konsonantencluster gestört (mittlere Häufigkeit)	.957 (.274)	1.327 (.719)	.698 (.232)	1.087 (.537)
Konsonantencluster ungestört (hohe Häufigkeit)	1.132 (.307)	1.533 (.758)	.778 (.183)	1.258 (.570)
Konsonantencluster gestört (hohe Häufigkeit)	.934 (.208)	1.384 (.648)	.628 (.262)	.857 (.269)

Tabelle 2: Durchschnittliche Latenzzeiten der Pseudowörter nach Gruppen (in Sekunden) pro Schulstufe, aufgeteilt in gute/schlechte LeserInnen. Standardabweichungen in Klammern.

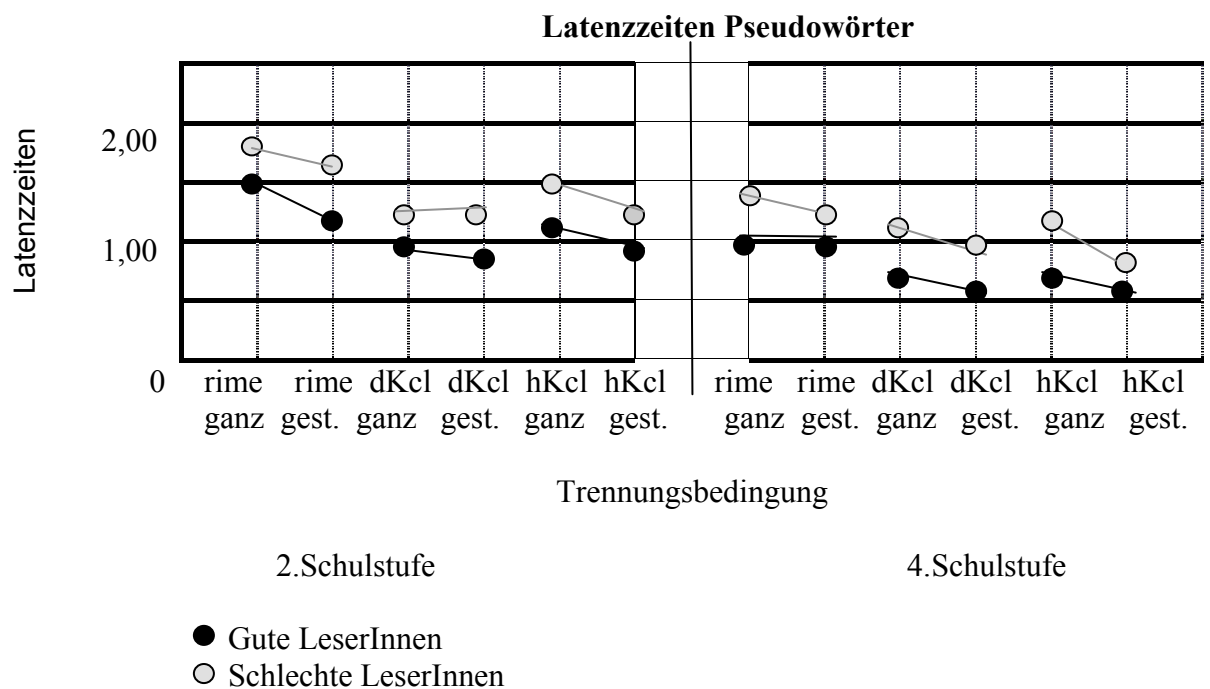


Abbildung 6: Durchschnittliche Latenzzeiten der Pseudowörter pro Schulstufe in Sekunden, aufgeteilt in gute/schlechte LeserInnen.

Es gab einen Haupteffekt des Zwischenfaktors Gruppe, $F(3, 84) = 7.22$, $p = .000$, $r = .205$. Jene SchülerInnen, die in der Vorerhebung langsamer lasen, erzielten auch in der Haupterhebung höhere Latenzzeiten.

Bei beiden Zwischensubjektfaktoren konnte ein Haupteffekt der Störung beobachtet werden, $F(1, 84) = 79.58$, $p = .000$, $r = .486$ für Gruppe und $F(1, 86) = 73.98$, $p = .000$, $r = .462$ für Schulstufe. Die angeforderten Kontraste zeigten beim Faktor Gruppe, $F(1, 84) = 79.58$, $p = .000$, $r = .486$, und beim Faktor Schulstufe, $F(1, 86) = 73.98$, $p = .000$, $r = .462$, dass Pseudowörter mit ungestörtem Cluster langsamer erlesen wurden, als jene mit gestörtem Cluster.

Es gab weiters einen Interaktionseffekt zwischen Störung und Gruppe, $F(3, 84) = 2.91$, $p = .039$, $r = .094$. Kontraste bestätigten, dass gestörte Cluster von guten LeserInnen der 2.Klassen schneller benannt wurden als von schlechten LeserInnen der 4.Klassen.

Analysen der einzelnen Gruppen

2. Schulstufe

Gute LeserInnen

Es zeigte sich ein genereller Haupteffekt der Störung, $F(1, 21) = 31.21$, $p = .000$, $r = .598$. Die angeforderten Kontraste bestätigten, dass Pseudowörter mit ungestörtem Cluster immer langsamer erlesen wurden, als jene mit gestörtem Cluster.

Schlechte LeserInnen

Auch in dieser Gruppe zeigte sich ein Haupteffekt der Störung, $F(1, 21) = 10.85$, $p = .003$, $r = .341$. Kontraste bestätigten, dass Pseudowörter mit ungestörtem Cluster immer höhere Latenzzeiten aufwiesen, als jene mit gestörtem Cluster.

4. Schulstufe

Gute LeserInnen

Es zeigte sich ein übergreifender Haupteffekt der Störung, $F(1, 21) = 18.257$, $p = .000$, $r = .465$. Die angeforderten Kontraste bestätigten, dass Pseudowörter mit ungestörtem Cluster langsamer gelesen wurden, als Pseudowörter mit gestörtem Cluster.

Analyse beider Gruppen

Konsonantencluster

Bei den Konsonantenclustern hoher Häufigkeit konnte ein Interaktionseffekt zwischen dem Cluster und der Gruppe beobachtet werden, $F(1, 42) = 7.51$, $p = .009$, $r = .152$. Kontraste zur Prüfung der Signifikanz zeigten, dass gute LeserInnen Pseudowörter mit gestörtem Cluster signifikant schneller lasen als schlechte LeserInnen dies taten.

Schlechte LeserInnen

Es konnte auch in dieser Gruppe ein Haupteffekt der Störung beobachtet werden, $F(1, 21) = 24.757$, $p = .000$, $r = .541$. Die angeforderten Kontraste zeigten, dass Pseudowörter mit ungestörtem Cluster höhere Latenzzeiten hatten als dies bei Pseudowörtern mit gestörten Clustern der Fall war.

Darüber hinaus konnte bei den Konsonantencluster mit hoher Häufigkeit ein Interaktionseffekt zwischen dem Cluster und der Gruppe beobachtet werden, $F(1, 42) = 7.51$, $p = .009$, $r = .152$. Kontraste zur Prüfung der Signifikanz zeigten, dass gute LeserInnen Pseudowörter mit gestörtem Cluster signifikant schneller lasen als schlechte LeserInnen.

Interaktionseffekte von Wörtern und Pseudowörtern

Da die Varianzanalysen bei Wörtern und Pseudowörtern zu gegensätzlichen Ergebnissen in den einzelnen Trennungsbedingungen führten, wurden letztlich auch die Interaktionseffekte zwischen Wörtern und Pseudowörtern der jeweiligen Clusterarten, getrennt nach Schulstufe, berechnet.

2. Schulstufe

Sowohl bei Konsonantencluster durchschnittlicher, als auch bei jenen mit hoher Häufigkeit, konnten Interaktionseffekte beobachtet werden.

Wörter mit ungestörtem Konsonantencluster durchschnittlicher Häufigkeit wurden signifikant schneller erlesen als Wörter mit gestörtem Cluster $F(1, 42) = 13.72$, $p = .001$, $r = .246$. Dieser Effekt trat auch bei Wörtern mit Konsonantencluster hoher Häufigkeit auf, $F(1, 42) = 20.23$, $p = .000$, $r = .325$.

Bei den Pseudowörtern zeigte sich ein gegenteiliges Ergebnis. Diejenigen mit gestörtem Konsonantencluster wurden signifikant schneller erlesen, als jene mit ungestörtem Cluster.

4. Schulstufe

Es zeigten sich signifikante Interaktionseffekte in der Onset/Rime Bedingung und bei den Konsonantenclustern hoher Häufigkeit.

Wörter mit ungestörtem Onset/Rime Cluster wurden schneller gelesen als Wörter mit gestörtem Cluster, $F(1, 42) = 5.71$, $p = .021$, $r = .120$. Derselbe Effekt zeigte sich bei den Konsonantenclustern hoher Häufigkeit, $F(1, 42) = 30.74$, $p = .000$, $r = .423$. Bei den Pseudowörtern zeigte sich ein gegenteiliges Ergebnis. Diejenigen mit gestörtem Cluster wurden nämlich schneller erlesen, als Pseudowörter mit ungestörtem Cluster.

5. Diskussion

5.1 Interpretation

Das Hauptergebnis der vorliegenden Studie war, dass sich die durchschnittlichen Latenzzeiten der einzelnen Trennungsbedingungen teilweise signifikant voneinander unterschieden, dieser Effekt jedoch nicht generalisiert werden kann.

Die Analysen der einzelnen Gruppen pro Schulstufe zeigten, dass nur bei den schlechten LeserInnen der 2. Schulstufe ein Haupteffekt der Störung beobachtet werden konnte. Wörter mit gestörtem Cluster wurden demnach langsamer erlesen, als jene mit ungestörtem Cluster. Bei allen anderen Gruppen zeigte sich dieser Effekt nicht. Die Verarbeitung sublexikaler Einheiten scheint aus diesem Grund nur bei langsamen LeseanfängerInnen (2. Schulstufe) von der Lesekompetenz abzuhängen. Der angenommene Trennungseffekt auf Ebene der Schulstufen, welcher erst bei LeserInnen der vierten Klassen erwartet wurde, konnte in diesem Fall nicht bestätigt werden.

Um trotz allem verwertbare Rückschlüsse aus der Untersuchung ziehen zu können, wurden darüber hinaus die Wörter der einzelnen Clusterarten, aufgeteilt nach Schulstufe, analysiert. Hier zeigte sich, dass, außer in der Onset/Rime Trennungsbedingung in der 2. Schulstufe, ein

Haupteffekt der Störung beobachtet werden konnte. Ein bemerkenswerter Interaktionseffekt zeigte sich dabei bei den durchschnittlichen Konsonantenclustern in der vierten Schulstufe. Hierbei wurden nämlich Wörter mit einem gestörten Cluster schneller erlesen als Wörter mit einem ungestörten Cluster. Ansonsten wiesen Wörter mit einem gestörten Cluster höhere Latenzzeiten auf, als diejenigen mit einem ungestörten Cluster. Erwähnenswert ist auch, dass alle Wörter und Pseudowörter der Onset/Rime Bedingungen in allen Gruppen die höchsten Latenzzeiten aufwiesen und demnach am langsamsten erlesen wurden.

Ein weiteres interessantes Ergebnis war, dass Wörter und Pseudowörter gegensätzliche Ergebnisse lieferten. Die durchschnittlichen Latenzzeiten in Tabelle 2 zeigen, dass Pseudowörter mit gestörtem Cluster in allen Trennungsbedingungen schneller erlesen wurden, als diejenigen mit ungestörtem Cluster. Diese Tatsache führte auch zu einigen signifikanten Interaktionseffekten.

Eine mögliche Erklärung dafür könnte sein, dass Wörter und deren dazugehöriges Buchstabencluster bereits im mentalen Gedächtnis gespeichert wurden. Eine Störung des verinnerlichten Clusters erhöht in diesem Zusammenhang die Latenzzeit der Wortbenennung. Es ist jedoch nicht einfach ein schlüssiges Fazit aus den gegensätzlichen Ergebnissen der Pseudowörter zu ziehen. Eine Störung von sublexikalischen Clustern scheint in diesem Fall einen positiven Effekt auf den Lesefluss zu haben. Beim Vorlesen von Pseudowörtern müssen Kinder ein, für sie unbekanntes, Wort selbst strukturieren. Eine Störung von eigentlich bekannten Einheiten scheint hierbei als Hilfestellung zu dienen. Dieser Befund steht somit im Widerspruch zu den Ergebnissen der Wörter und den Theorien des Lesens. Eine mögliche Ursache könnte in den Lauten liegen, die den Clustern folgen. Bei den Pseudowörtern wurden oftmals Verschlusslaute („b“ oder „p“) nach den betroffenen Clustern verwendet, was durchaus einen Einfluss auf das Lesen haben könnte. Ein enger Zusammenhang zur phonologischen Bewusstheit, welche auch von Klicpera et al. (2001) betont wird, kann somit angenommen werden.

Eine weitere Erklärung könnte der sogenannte Positionseffekt sein. Da die eingebauten Störungen im hinteren Teil des Wortes eingebaut wurden, könnte dies so einen Einfluss auf den Lesevorgang haben.

Bezugnehmend auf die beiden Zwischensubjektfaktoren Gruppe (gute/schlechte LeserInnen) und Schulstufe, zeigt sich, dass jene SchülerInnen, die in der Vorerhebung als langsame LeserInnen identifiziert wurden, auch in der Haupterhebung höhere Latenzzeiten aufwiesen. Dieselbe Schlussfolgerung lässt sich auch bei den guten LeserInnen ziehen. Das Salzburger Lesescreening stellte demnach eine gute Methode dar, generell zwischen guten und schlechten LeserInnen zu differenzieren.

Weiters lässt sich sagen, dass die Leseleistung der SchülerInnen mit fortschreitendem Leseunterricht grundsätzlich steigt. Gute LeserInnen der 4. Schulstufe erreichten durchschnittlich geringere Latenzzeiten als gute LeserInnen der 2. Schulstufe. Wie bereits in Punkt 4 angeführt, zeigten paarweise Vergleiche, dass es keinen signifikanten Unterschied zwischen den guten LeserInnen der 2. Schulstufe und den schlechten LeserInnen der 4. Schulstufe gab.

Zudem konnte bei den Pseudowörtern beobachtet werden, dass gute LeserInnen der 2. Klassen höhere durchschnittliche Latenzzeiten erreichten als schlechte LeserInnen der 4. Klassen.

Wie zuvor erwähnt, stellten Wörter mit gestörtem Konsonantencluster durchschnittlicher Häufigkeit für gute als auch für schlechte LeserInnen in der 4. Schulstufe keine Beeinträchtigung des Leseflusses dar. Das könnte bedeuten, dass diese nicht als sublexikale Einheiten im mentalen Lexikon gespeichert sind, oder sich fortgeschrittene LeserInnen durch eine visuelle Störung wie diese nicht mehr beeinflussen lassen.

In dieser Studie wurde eine Bedingung erstellt, die sich auf Konsonantencluster mit einer hohen Häufigkeit konzentrierte und welche auch im schulischen Kontext bereits als Einheit erlernt werden. Die Mittelwerte in Tabelle 1 zeigen, dass Wörter mit gestörtem Konsonantencluster einer hohen Häufigkeit, außer in der Gruppe der schlechten LeserInnen der 2. Schulstufe, langsamer erlesen wurden als dies bei Wörtern mit ungestörtem Cluster der Fall war. Diese Ergebnisse sprechen für eine Verwendung dieser sublexikalischen Einheiten beim Lesen. Da die schlechten LeserInnen der 2. Klassen diesen Resultaten widersprechen, kann davon ausgegangen werden, dass sie Wörter mit diesen Clustern noch nicht verinnerlicht und gespeichert haben.

Bei den Pseudowörtern ergibt sich in diesem Punkt ein ähnliches Bild. Mit Ausnahme der schlechten LeserInnen der 2. Schulstufe erreichten Pseudowörter mit gestörtem

Konsonantencluster hoher Häufigkeit bei allen Gruppen niedrigere Latenzzeiten als bei Pseudowörtern durchschnittlicher Häufigkeit. Dies lässt nun den Schluss zu, dass nur richtige Wörter mit ihren sublexikalischen Teilen im mentalen Lexikon gespeichert sind und eine Störung den Leseprozess beeinträchtigt.

Hauptziel dieser Diplomarbeit war zu untersuchen, inwieweit sublexikale Cluster beim Lesen als Einheit verarbeitet werden und ob der Lesefluss durch eine visuelle Segmentierung dieser Cluster beeinflusst werden kann. Des Weiteren sollte auch untersucht werden, ob ein Zusammenhang mit dem Lesefortschritt (AnfängerInnen/fortgeschrittene LeserInnen) der Kinder besteht.

Im Folgenden sollen nun die oben angeführten Ergebnisse mit den in Kapitel 2 erwähnten Entwicklungsmodellen und Lesetheorien in Verbindung gebracht werden.

Nach Ehris Phasentheorie (1995) kann erst in der fortgeschrittenen konsolidierten alphabetischen Phase von einer Automatisierung des Leseverhaltens ausgegangen werden. Dabei werden Silben, Onsets oder auch Rimes als Einheiten im Gedächtnis gespeichert und das Lesen dadurch vereinfacht, wobei Wörter noch nicht holistisch erkannt werden. Nach Ehris (1995) befinden sich erst LeserInnen in einem fortgeschrittenen Stadium in dieser Phase und können von dieser Strategie profitieren. Diesen Annahmen nach sollten fortgeschrittene LeserInnen, in diesem Fall jene der vierten Schulstufe, einen Trennungseffekt aufweisen und in ihrem Lesefluss beeinträchtigt werden.

Diese Ansichten stehen im engen Zusammenhang mit den Annahmen von Klicpera et al. (2010), die ebenfalls erst in der fortgeschrittenen alphabetischen Phase mit voller Integration von einer zunehmenden Automatisierung des Leseprozesses ausgehen. In diesem Zusammenhang werden nach diesem Modell mentale Einheiten gebildet, die häufig vorkommende Buchstabencluster (wie beispielsweise „ck“ oder „ch“) beinhalten und somit das Lesen erleichtern. Bezogen auf die vorliegende Studie würden sich, nach Ehris Phasentheorie (1995) und dem Kompetenzentwicklungsmodell von Klicpera et al. (2001), bereits gute LeserInnen der zweiten Schulstufe in diesem Stadium befinden.

Auch Frith (1985) spricht in seiner Theorie des Leselerwerbs mit dem Erreichen der orthografischen Phase von einer Zerlegung der Wörter in kleinere orthografische Einheiten und somit für die Speicherung sublexikalischer Cluster im Gedächtnis.

Schließlich geht auch die Grain Size Theorie von Ziegler und Goswami (2005) von einer Verwendung sublexikaler Einheiten („grain sizes“) beim Lesen aus. Die Autoren betonen hierbei weiters die Wichtigkeit der Konsistenz, beziehungsweise Inkonsistenz einer Sprache. Demnach scheint die Abspeicherung von sprachlichen Einheiten bei inkonsistenten Sprachen eine zunehmende Erleichterung und Beschleunigung des Leseprozesses zu sein. Deutsch ist zwar eine relativ konsistente Sprache, trotz allem scheinen nach Ziegler und Goswami (2005) LeserInnen größere sublexikale Cluster zu speichern, um das Lesen zu vereinfachen.

Bezugnehmend auf die Modelle des kompetenten Lesens kann gesagt werden, dass sich sowohl das Parallel Distributed Modell von Seidenberg und McClelland (1989) und auch das Konnektionistische Zwei-Wege-Modell (CDP++) von Perry, Ziegler und Zorzi (2010) für die Berücksichtigung sublexikaler Einheiten während des Lesens aussprechen und somit die Ergebnisse der vorliegenden Studie unterstreichen. Bei beiden scheinen jene Einheiten ein fester Bestandteil des Modells zu sein, welche demnach auch im Leseprozess berücksichtigt werden.

Lediglich das DRC Modell von Coltheart et al. (2001) spricht gegen eine Verwendung sublexikaler Cluster beim Lesen und somit auch gegen die Ergebnisse der Wörter in der vorliegenden Untersuchung. In diesem Modell werden Einheiten dann berücksichtigt, wenn die Buchstabengruppe einen Laut repräsentiert (beispielsweise „sp“ oder „st“), diese werden jedoch nicht für eine weitere Verwendung im Modell als Einheit gespeichert. Diese Annahmen wiederum lassen sich mit den Ergebnissen der Pseudowörter in Einklang bringen.

Die Schlussfolgerungen, dass sich die Latenzzeiten tendenziell eher erhöhen, wenn eine visuelle Störung im Rime eingebaut wird, als wenn sich diese zwischen Onset und Rime befindet, stimmen im Großen und Ganzen mit den Ergebnissen von Treiman und Chafetz (1987) überein. Auch Bowey (1990) gelangte mit ihren Priming Experimenten zu demselben Schluss. Eine weitere Untersuchungsreihe von Bowey (1990) verwendete dasselbe visuelle Segmentierungsparadigma wie die vorliegende Untersuchung. Hier gelang man ebenfalls zu der Schlussfolgerung, dass Onset und Rime als Einheiten beim Lesen verarbeitet werden und bestätigt demnach die Ergebnisse der Studie dieser Diplomarbeit.

Levitt, Healy und Fendrich (1991) kamen zu gegenteiligen Resultaten. Bei ihren Experimenten konnte keine Gliederung des Wortes in Onset und Rime nachgewiesen werden. Zu demselben Schluss kamen auch Marinus und de Jong (2008) in ihrer Untersuchungsreihe,

welche demnach nicht von einer Verarbeitung von Onset und Rime als Einheit ausgehen. Auch Mousikou et al. (2010) widersprechen der Annahme, dass Onsets als Cluster verarbeitet werden.

Da sich die in der vorliegenden Studie verwendeten Konsonantencluster am Ende des Wortes befanden, konnten die Ergebnisse von Frederiksen und Kroll (1976) nicht überprüft werden. Die Autoren fanden nämlich heraus, dass Konsonantencluster am Anfang eines Wortes höhere Latenzzeiten verursachen als am Ende.

Bezugnehmend auf die Verwendung von Konsonantencluster mit hoher und niedriger Häufigkeit, lässt sich eine Verbindung zur Studie von Hennessey und Kirsner (1999) herstellen. Diese fanden heraus, dass sublexikale Prozesse bei Wörtern mit niedriger Häufigkeit vermehrt involviert sind, als dies bei Wörtern mit hoher Häufigkeit zutrifft. Die Ergebnisse der durchgeführten Studie dieser Diplomarbeit zeigen, dass alle Gruppen, bis auf die schlechten LeserInnen der 2. Schulstufe, Wörter mit gestörtem Konsonantencluster mittlerer Häufigkeit schneller benennen als dies bei Konsonantencluster von hoher Häufigkeit der Fall ist. Eine Störung eines bekannten Konsonantenclusters führt demnach zu verlangsamten Latenzzeiten. Diese Ergebnisse widersprechen im Grunde jenen von Hennessey und Kirsner (1999).

Zusammenfassend kann nun gesagt werden, dass Modelle des Leseerwerbs im Grunde davon ausgehen, dass sublexikale Cluster als solche gespeichert und auch verarbeitet werden. Denselben Standpunkt vertreten auch, bis auf das DRC Modell von Coltheart et al. (2001), Theorien des kompetenten Lesens. Die durchgeführte Studie kann hinsichtlich der Verarbeitung von Wörtern demnach im Großen und Ganzen mit dem dargestellten theoretischen Hintergrund in Einklang gebracht werden. Ein inkonsistenter und deshalb klärungsbedürftiger Befund dabei ist, dass Wörter und Pseudowörter gegenteilige Ergebnisse aufweisen und die Ergebnisse der letztgenannten Modelle widersprechen und nur den Annahmen des DRC Modells entsprechen.

Es scheint, als ob auch bei schlechten LeserInnen der 2. Schulstufe Wörter und ihre sublexikalischen Einheiten als solche bereits im mentalen Lexikon abgespeichert sind und eine visuelle Störung den Lesefluss beeinträchtigt.

Da in der Vergangenheit keine eindeutigen Forschungsergebnisse zu diesem Themengebiet gefunden werden konnten, scheint es umso interessanter zu sein, weiter vertiefende Untersuchungen in diese Richtung durchzuführen und dadurch weitere neue Einblicke zu erlangen.

5.2 Kritische Diskussion

Wie bereits erwähnt, gibt es bislang nur wenige Studien, die sich mit der Verarbeitung von sublexikalischen Clustern, vor allem auch mit diesem speziellen Segmentierungsparadigma, beschäftigen. Da es aber durchaus eine wesentliche Bedeutung für das Leseverhalten zu haben scheint, werden im Folgenden einzelne Aspekte der vorliegenden Studie kritisch diskutiert, um diese gegebenenfalls in weiterführenden Untersuchungen zu optimieren.

Stichprobengröße

Die Studie wurde zwar mit 232 SchülerInnen geplant, was eine ausreichende Anzahl für dieses Projekt darstellen würde, letztendlich nahmen jedoch nur 179 Kinder aufgrund mangelnder Einverständniserklärungen der Eltern daran teil. Darüber hinaus mussten zwölf weitere Kinder wegen zu niedriger Leistungen in der Vorerhebung ausgeschlossen werden, wodurch sich die Stichprobengröße abermals auf 166 SchülerInnen verringerte.

Um noch aussagekräftigere Ergebnisse zu erlangen, sollte in zukünftigen Arbeiten auf eine größere TeilnehmerInnenanzahl geachtet werden. Um mangelnde Einverständniserklärungen, welche im Pflichtschulbereich unabdingbar sind, vorzubeugen, sollte eine noch größere Anzahl an Elternbriefen verteilt, beziehungsweise mehr Schulen rekrutiert werden, um einem Schwund wie diesem vorzubeugen.

Silbenhäufigkeiten

In der Untersuchung wurden nicht nur Cluster durchschnittlicher Häufigkeit erforscht, sondern es wurde auch explizit eine Bedingung mit oft vorkommenden Konsonantenclustern erstellt.

Durch diese Vorgehensweise konnte auch untersucht werden, ob sublexikale Einheiten, wie „st“ oder „ck“, welche im schulischen Kontext in den meisten Fällen bereits als Cluster erlernt werden, auch als solche beim Lesen verarbeitet werden und eine visuelle Störung den Lesefluss beeinträchtigt.

Erstellung der Pseudowörter

Aufgrund der gefundenen Ergebnisse könnte bei der Auswahl der Pseudowörter verstärkt darauf geachtet werden, welche Phoneme den jeweiligen sublexikalischen Clustern folgen. Dadurch könnte auch ein Zusammenhang mit der phonologischen Bewusstheit untersucht werden.

Segmentierungsparadigma

Nach Wissen der Autorin wurden bislang nur wenige Studien durchgeführt, die mit einer visuellen Segmentierung in Form von alternierenden Groß- und Kleinbuchstaben arbeiteten. Da Hauptwörter in der deutschen Sprache nur zu Beginn großgeschrieben werden, könnte der Wechsel im Wort eine zunehmende Beeinträchtigung des Leseflusses darstellen, als dies beispielsweise eine Trennung in Form eines Symbols (# oder //) wie bei Marinus und de Jong (2008) und Martensen et al. (2003) darstellen würde. In einer breiter angelegten Studie könnten auch beide Segmentierungsarten verwendet werden um herauszufinden, ob sich hierbei ein Unterschied beobachten lässt und ob sich die hier gefundenen Effekte auch in einer weiteren Trennungsart zeigen.

Trennungsbedingungen

Die visuelle Trennung wurde in der Gruppe der Konsonantenclustern entweder so verwendet, dass es die Einheit störte, oder dass sie als solche erhalten blieb. Jene Segmentierung befand sich hierbei stets im hinteren Teil des Wortes. Da es sich bei der visuellen Störung um alternierende Groß- und Kleinbuchstaben handelte, hätte eine Trennung im vorderen Bereich zur Folge, dass ein reguläres deutsches Wort entstünde (zum Beispiel: Geld) und somit von keiner Beeinträchtigung des Leseflusses in diesem Sinne gesprochen werden könnte. In weiteren Studien könnte auch eine weitere Bedingung eingeführt werden, in der das Wort entweder nur in Klein- oder Großbuchstaben präsentiert wird, um die durchschnittlichen Latenzzeiten dieser Kontrollbedingung mit jenen der gestörten/ungestörten Bedingung zu vergleichen.

5.3 Ausblick und praktische Relevanz

Es gibt bislang nur eine geringe Anzahl vergleichbarer Versuchsdesigns, die sich mit der funktionellen Verarbeitung sublexikaler Cluster während des Leseprozesses beschäftigen. Die durchgeführte Studie sollte herausfinden, ob diese Einheiten als solche beim Lesen verarbeitet werden.

Aufgrund der Ergebnisse kann davon ausgegangen werden, dass sublexikale Cluster durchaus als Einheit verarbeitet und gespeichert werden. Bezugnehmend auf den Lesefortschritt zeigte sich, dass sich vor allem „schlechte LeseanfängerInnen“ (schlechte LeserInnen der zweiten Schulstufe) durch eine Trennung des Clusters beeinflussen lassen.

Es gilt durchaus als gängige Methode im Unterricht, SchülerInnen auf eine silbenweise Trennung beziehungsweise Betonung der Wörter hinzuweisen. Da Onsets und Rimes im Grunde kleinere Einheiten von Silben sind, kann es demnach auch von Vorteil sein, im schulischen Kontext in Zukunft auch auf diese Cluster zu achten und diese zu betonen.

Die Ergebnisse zeigen weiters, dass es sich zu bewähren scheint, Konsonantencluster, welche auch ein Phonem bilden (zum Beispiel „ck“ oder „st“), bereits im Erstleseunterricht als Einheit zu erarbeiten sodass diese in weiterer Folge als solche von den Kindern gespeichert werden können. All jene Aspekte können den Lesefluss positiv beeinflussen und leseschwachen SchülerInnen eine wertvolle Hilfestellung bieten.

Da bei den Pseudowörtern jene mit gestörtem Cluster schneller erlesen wurden, kann hier eine enge Verbindung zur phonologischen Bewusstheit angenommen werden. Diese Tatsache stellt eine weitere interessante Forschungsmöglichkeit dar,

Theoretisch bedeutsam wäre es auch herauszufinden, inwieweit sich legasthene Kinder von den Störungsbedingungen, wie sie in dieser Studie angewendet wurden, beeinflussen lassen. Fördermaterialien und gezielte Lesetrainings könnten hierbei auf alle Fälle erfolgreich angewendet werden.

Das Arbeiten mit sublexikalischen Clustern könnte demnach in den Schulbüchern ein fester Bestandteil werden und somit in den Leseunterricht integriert werden, um SchülerInnen individuell fördern zu können.

In welchem Ausmaß Materialien wie diese eingesetzt werden sollten und wie diese zu gestalten sind, um einen positiven Effekt auf den Leseprozess zu erzielen, bedarf jedoch weiterer Diskussion und würde den Rahmen dieser Diplomarbeit überschreiten.

6. Zusammenfassung

6.1 Deutsche Zusammenfassung

Entwicklungsmodelle des Lesens, wie beispielsweise jenes von Ehri (1995) oder Frith (1985), gehen davon aus, dass die Verarbeitung und der Gebrauch sublexikaler Cluster erlernt werden und dass diese Verknüpfungen im Lesesystem fest etabliert sind.

Das Netzwerkmodell von Seidenberg und McClelland (1989) oder das CDP++ Modell von Perry et al. (2010), welche sich mit dem Ablauf des Leseprozesses beschäftigen, behaupten, dass sublexikale Cluster gespeichert und immer wieder darauf zurückgegriffen werden kann. Lediglich das DRC Modell von Coltheart et al. (2001) widersprach der Berücksichtigung der sublexikalischen Einheiten beim Lesen.

In der vorliegenden Studie wurde die Verarbeitung sublexikalischer Einheiten von SchülerInnen der 2. und 4. Schulstufe untersucht. Diese Diplomarbeit konzentrierte sich dabei einerseits auf den Onset und den Rime eines Wortes beziehungsweise Pseudowortes und andererseits auf Konsonantencluster. Im zweiten Fall wurden nicht nur Einheiten mit einer durchschnittlichen, sondern auch mit einer hohen Häufigkeit untersucht.

Die sublexikalischen Cluster wurden dabei durch ein visuelles Segmentierungsparadigma, bestehend aus alternierenden Groß- und Kleinbuchstaben (zum Beispiel: gELD), getrennt. Dabei wurde die Hypothese aufgestellt, dass jene Wörter und Pseudowörter, deren Cluster gestört sind, langsamer gelesen werden und im weiteren Sinne höhere Latenzzeiten aufweisen, als dies bei Wörtern und Pseudowörtern mit ungestörtem Cluster der Fall ist.

In der Vorerhebung wurde die Leseleistung von insgesamt 178 SchülerInnen der 2. und 4. Schulstufe durch das Salzburger Lesescreening (Mayringer & Wimmer, 2003) erfasst. Um auch die Auswirkungen auf die unterschiedlichen Lesekompetenzen zu untersuchen, wurden die 22 schnellsten und langsamsten LeserInnen pro Schulstufe ermittelt und in die Haupterhebung aufgenommen. Hier sollten die Kinder vorgegebene Wörter und Pseudowörter laut vorlesen. Sie wurden dabei mit einem Audioprogramm aufgenommen, sodass die einzelnen Latenzzeiten später ermittelt und kodiert werden konnten. Des Weiteren wurden nur richtig erlesene Wörter und Pseudowörter in die Auswertung mit einbezogen.

Die Ergebnisse zeigen, dass sich die durchschnittlichen Latenzzeiten der einzelnen Trennungsbedingungen meist signifikant voneinander unterscheiden. Eine Analyse der

einzelnen Gruppen pro Schulstufe veranschaulichte, dass lediglich bei den schlechten LeserInnen der 2. Schulstufe ein Haupteffekt der Störung zu beobachten war. Wörter mit gestörtem Cluster wurden demnach langsamer erlesen als Wörter mit ungestörtem Cluster. Diese Tatsache lässt den Schluss zu, dass Kinder mit dieser Lesekompetenz sehr anfällig für eine Störung bereits gespeicherter Einheiten sind.

Eine weitere Analyse der einzelnen Clusterarten pro Schulstufe zeigte, dass, außer in der Onset/Rime Trennungsbedingung in der 2. Schulstufe, ein Haupteffekt der Störung auf die Latenzzeit verzeichnet werden konnte. Alle anderen Wörter mit einem gestörten Cluster wurden langsamer erlesen, als jene mit einem ungestörten Cluster. Eine Ausnahme bildeten hier Wörter mit einem durchschnittlichen Konsonantencluster: Wörter mit einem gestörten Cluster wiesen in dieser Bedingung höhere Latenzzeiten auf, als diejenigen mit einem ungestörten Cluster. Diese Einheiten scheinen nicht als solche im mentalen Gedächtnis gespeichert zu sein.

Auffällig war, dass bei Wörtern und Pseudowörtern gegensätzliche Ergebnisse beobachtet werden konnten. Pseudowörter mit einem gestörten Cluster wurden nämlich schneller erlesen als mit einem ungestörten Cluster. Eine Erklärung hierfür könnte in einem Positionseffekt liegen, da die Störung im hinteren Teil des Wortes eingebaut war. Eine weitere mögliche Begründung könnte sein, dass eine Störung der Einheit in diesem Fall einen positiven Effekt auf den Leseprozess hat.

Die Ergebnisse der Wörter stehen im Großen und Ganzen im Einklang mit den angeführten Leseerwerbstheorien und gängigen Theorien des kompetenten Lesens. Jene des DRC Modells von Coltheart (2001) widersprechen ihnen, stehen jedoch im Einklang mit den Ergebnissen der Pseudowörter. Eine mögliche Erklärung könnte sein, dass das schnellere Lesen durch die Phoneme, welche den betroffenen Clustern folgen, beeinflusst wird.

Individuelle Fördermaßnahmen scheinen aufgrund der Ergebnisse praktisch wertvoll zu sein. In welchem Ausmaß dies umzusetzen ist und die Frage, ob die Ergebnisse der Pseudowörter eventuell mit der Entwicklung der phonologischen Bewusstheit zusammenhängen, bedarf jedoch weiterer Diskussion und Forschung, welche jedoch den Rahmen dieser Arbeit überschreiten würden.

Die vorliegenden Resultate zeigen jedoch, dass darauf ausgerichtete Maßnahmen, vor allem für SchülerInnen in der 2. Schulstufe, welche eine schlecht ausgeprägte Lesekompetenz besitzen, durchaus von praktischem Nutzen sein kann.

6.2 Summary

Developmental theories, like the ones of Ehri (1995) or Frith (1985), affirm that the use of sublexical clusters can be learned and that these relationships are inherent parts of the human reading system.

The connectionist models of Seidenberg and McClelland (1989) and the CDP++ model of Perry et al. (2010) claim, that these clusters are retained in the reading system and can be retrieved if needed. Simply the DRC model of Coltheart et al. (2001) is contrary to these assumptions.

In this study the use of sublexical letter clusters of children in second and fourth grade in elementary school was examined. The research focused on the one hand on the onset and rime of a word and on the other on consonant clusters. In reference to the consonants, not only items with an average frequency but also those with a high frequency were investigated. The study used a segmentation paradigm by separating two adjacent letters with alternating upper and lower case letters (for example: gELD).

It was hypothesized that those words, whose cluster has been separated, are read more slowly than those words whose cluster is intact.

In a pre-investigation the reading skills of 178 pupils of second and fourth grade of elementary school were tested. After this screening the 22 fastest and slowest readers of each grade were identified and affiliated to the main-investigation where the pupils had to read out words and nonwords. They were recorded with a special audio programme so that latencies could be determined afterwards. Beyond that, only words and nonwords that have been named correctly were included in further evaluation.

The results showed that the mean latencies in most of the conditions are significantly different. An analysis of the different groups (slow/fast reader) showed that only slow readers of second grade read words with separated cluster more slowly than those with intact cluster. It seems that pupils in this reading stage are more prone to a separation of already internalized letter clusters.

A further analysis of the particular types of cluster has been made. This showed that, except words of the Onset/Rime condition, words with separated sublexical cluster were read significantly more slowly than words with intact cluster. Though there has been one exception: words that had a mean frequency consonant cluster were read faster if the letter

cluster was separated. It seems as if sublexical clusters like this are not stored in the mental lexicon.

A very noticeable outcome was furthermore that words and nonwords showed oppositional results: nonwords with a separated sublexical cluster were read faster than nonwords whose cluster was intact. This could be illustrated by a position effect of the segmentation paradigm. Another possible explanation could be that the segmentation paradigm has a positive effect on the reading speed.

The presented outcomes are consistent with developmental theories and theories of skilled reading in most instances. Merely the DRC model and the outcomes of words are contradictory. But otherwise they are in common with the findings of nonwords. One possible explanation could be that the phonemes that follow the letter cluster have a certain influence on the pupils' reading fluency. This could be an indication of a relationship to the development of the phonological awareness of children.

It seems as if special methods for fostering skilled reading could be very useful for slow readers in second grade of elementary school. For a detailed implementation further research is needed.

7. Literaturverzeichnis

- Baayen, R.H., Piepenbrock, R. & van Rijn, H. (1993). The CELEX Lexical Database [CD-ROM]. Philadelphia: Linguistic Data Consortium, University of Pennsylvania.
- Baron, J. & Strawson, C. (1976). Use of orthographic and word-specific knowledge in reading words aloud. *Journal of Experimental Psychology: Human Perception and Performance*, 2, 386-393.
- Bowey, J. A. (1990). Orthographic onsets and rimes as functional units of reading. *Memory & Cognition*, 18, 419-427.
- Bowey, J. A. (1996). Orthographic onsets as functional units of adult word recognition. *Journal of Psycholinguistic Research*, 25, 571-595.
- Bruck, M., Caravolas, M. & Treiman, R. (1995). Role of syllable in the processing of spoken English: Evidence from a nonword comparison task. *Journal of Experimental Psychology*, 21, 469-479.
- Coltheart, M., Rastle, K., Perry, C., Langdon, R. & Ziegler, J. (2001). DRC: A Dual Route Cascaded Model of visual word recognition and reading aloud. *Psychological Review*, 108, 204-256.
- Ehri, L. C. (1995). Phases of development in learning to read words by sight. *Journal of Research in Reading*, 18, 116-125.
- Frederiksen, J. R. & Kroll, J. F. (1976). Spelling and sound: approaches to the internal lexicon. *Journal of Experimental Psychology*, 2, 361-379.
- Frith, U. (1985). Beneath the surface of developmental dyslexia. In K. E. Patterson, J. C. Marshall & M. Coltheart (Hrsg.), *Surface Dyslexia. Neuropsychological and Cognitive Studies of Phonological Reading* (S. 301-330). London: Erlbaum.
- Grainger, J. & Ferrand, L. (1996). Masked orthographic and phonological priming in visual word recognition and naming: cross-task comparisons. *Journal of Memory and Language*, 35, 623-647.
- Goswami, U. (1999). Causal connections in beginning reading: the importance of rhyme. *Journal of Research in Reading*, 22, 217-240.
- Hennessey, N. W. & Kirsner, K. (1999). The role of sub-lexical orthography in naming: a performance and acoustic analysis. *Acta Psychologica*, 103, 125-148.
- Klicpera, C., Schabmann, A. & Gasteiger-Klicpera, B. (2010). *Legasthenie – LRS. Modelle, Diagnosen, Therapie und Förderung*. München: Ernst Reinhardt Verlag.

- Levitt, A., Healy, A. F. & Fendrich, D. W. (1991). Syllable-internal structure and the sonority hierarchie: Differential evidence from lexical decision, naming and reading. *Journal of Psycholinguistic Research*, 20, 337-363.
- Marinus, E. & de Jong, P. F. (2008). The use of sublexical clusters in normal and dyslexic readers. *Scientific Studies of Reading*, 12, 253-280.
- Martensen, H., Maris, E. & Dijkstra, T. (2003). Phonological ambiguity and context sensitivity: On sublexical clustering in visual word recognition. *Journal of Memory and Language*, 49, 375-395.
- Mayringer, H. & Wimmer, H. (2003). Salzburger Lesescreening für die Klassenstufen 1-4. Bern: Hans Huber.
- Mousikou, P., Coltheart, M., Saunders, S. & Yen, L. (2010). Is the orthographic/phonological onset a single unit in reading aloud? *Journal of Experimental Psychology*, 36, 175-194.
- OECD (2005). Definition und Auswahl von Schlüsselkompetenzen: Zusammenfassung. Zugriff am 16.01.2013. Verfügbar unter: <http://www.oecd.org/pisa/35693281.pdf>.
- Perry, C., Ziegler, J. C. & Zorzi, M. (2010). Beyond single syllables: Large-scale modeling of reading aloud with the Connectionist Dual Process (CDP++) model. *Cognitive Psychology*, 61, 106-151.
- Plaut, D. C., McClelland, J. L., Seidenberg, M. S. & Patterson, K. (1996). Understanding normal and impaired word reading: Computational principles in quasi-regular domains. *Psychological Review*, 103, 56-115.
- Seidenberg, M. S. & McClelland, J. L. (1989). A distributed, developmental model of word recognition and naming. *Psychological Review*, 96, 523-568.
- Treiman, R. & Chafetz, J. (1987). Are there onset- and rime-like units in printed words? In M. Coltheart (Hrsg.), *Attention and Performance XII. The Psychology of reading* (S. 281-298). London: Erlbaum.
- Ziegler, J. C. & Goswami, U. (2005). Reading acquisition, developmental dyslexia, and skilled reading across languages: A psycholinguistic grain size theory. *Psychological Bulletin*, 131, 3-29.

8. Anhang

Anhang A

Liste verwendeter Wörter im Wortleasetest

Wörter		
Trennung zwischen Onset und Rime	keine Trennung der Konsonantencluster durchschnittlicher Häufigkeit	keine Trennung der Konsonantencluster hoher Häufigkeit
sEIL	heLD	roCK
rAUB	heLM	reST
bURG	fuND	baCH
hUND	geLB	zoPF
rAUS	kuLT	raND
Trennung im Rime	Trennung der Konsonantencluster durchschnittlicher Häufigkeit	Trennung der Konsonantencluster hoher Häufigkeit
seiT	welT	boCK
tauB	kurZ	misT
bunT	funK	tucH
huhN	senF	kopF
lauS	zelT	rahM
Pseudowörter		
Trennung zwischen Onset und Rime	keine Trennung der Konsonantencluster durchschnittlicher Häufigkeit	keine Trennung der Konsonantencluster hoher Häufigkeit
sEHL	weLB	nuCK
dRUB	veLZ	hiST
bURB	fuHL	guCH
hUNK	zeLM	wuPF
fRUS	kuNT	raMT
Trennung im Rime	Trennung der Konsonantencluster durchschnittlicher Häufigkeit	Trennung der Konsonantencluster hoher Häufigkeit
selD	zelN	racK
pauB	kuS	fusT
bupT	fulM	alcH
hurT	nelN	supF
sluS	serT	ralS

Anhang B

Lebenslauf

Julia Kast

geboren am 24.10.1083 in Wien

Ausbildung

seit Oktober 2006

Diplomstudium Psychologie (Universität Wien)

September 2003 – Juni 2006

Volksschullehramtsstudium (PÄDAK Wien;
Diplomprüfung: Juni 2006)

Oktober 2002 – Juni 2003

Diplomstudium Psychologie (Universität Wien)

Juni 2002

AHS-Reifeprüfung (Gymnasium Neusiedl/See)

Berufliche Tätigkeiten

seit Jänner 2012

Volksschullehrerin (1110 Wien)

September 2011 – Juni 2012

Praktikum in der Schulpsychologie-Bildungsberatung
(LSR für Burgenland)

Juli 2010 – Dezember 2011

Teilzeitkraft im Einzelhandel (Lacoste; FOC
Parndorf)

März 2007 – Juni 2010

Teilzeitkraft im Einzelhandel (Helly Hansen; FOC
Parndorf)