



universität
wien

DIPLOMARBEIT

Titel der Diplomarbeit

Die Verarbeitung von morphologischen Strukturen bei
geübten Lesern

Eine empirische Studie an Kindern am Ende der Grundschule

Verfasserin

Claudia Colley

Angestrebter akademischer Grad

Magistra der Naturwissenschaften (Mag. rer. nat.)

Wien, 2015

Studienkennzahl lt. Studienblatt: A 298

Studienrichtung lt. Studienblatt: Psychologie

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

DANKSAGUNG

Allen voran möchte ich Herrn Ao. Univ.-Prof. Dr. Alfred Schabmann für die Betreuung meiner Diplomarbeit danken, die er trotz langer Unterbrechung des Studiums meinerseits fortführte. Ich bedanke mich besonders für die menschliche und fachliche Unterstützung, die Bereitschaft zur Auseinandersetzung und für das entgegengebrachte Vertrauen selbstständig arbeiten zu können.

Großer Dank gilt den PädagogInnen und DirektorInnen der teilnehmenden Schulen und ganz besonders den SchülerInnen, die durch ihre herzliche Kooperation die vorliegende Arbeit überhaupt erst ermöglicht haben.

Dr. Reinhard Neumeier danke ich für den gewinnbringenden wissenschaftlichen Diskurs, indem seine Bereitschaft sich mit meinen oft weitreichenden Gedanken und Fragen auseinanderzusetzen ausgesprochen motivierend wirkte und seine Unterstützung im Fokussieren auf die zentrale Forschungsfrage den Prozess des Schreibens erleichterte und positiv beeinflusste.

Monika Kraus danke ich für ihre große Hilfe im Verfassen der englischen Zusammenfassung und Karin Sabo für ihre sofortige Bereitschaft die Arbeit Korrektur zu lesen.

Besonders bedanken möchte ich mich auch bei meiner Studienkollegin Barbara Sonnleitner. Die gemeinsame Entwicklung des Designs und Bewältigung der ersten Problemstellungen war durch eine gegenseitig wertschätzende, kollegiale und auch humorvolle Zusammenarbeit geprägt.

Während meines Studiums und insbesondere in der intensiven Schreibphase der Diplomarbeit wurde ich liebevoll von meiner Familie, allen voran meinen Eltern und meinem Bruder, sowie zahlreichen FreundInnen unterstützt und begleitet. Ich bedanke mich herzlich dafür und all ihre aufbauenden, motivierenden und stärkenden Worte.

Inhaltsverzeichnis

Inhaltsverzeichnis.....	5
Abbildungsverzeichnis.....	7
Tabellenverzeichnis.....	8
1 Einleitung.....	9
2 Theoretische Grundlagen.....	10
2.1 Entwicklungsmodelle des Schriftspracherwerbs.....	10
2.1.1 Phasenmodell des Sichtwortlesens.....	11
2.1.2 Kompetenzentwicklungsmodell des Lesens.....	13
2.1.3 Psycholinguistische Grain Size Theorie.....	15
2.2 Modelle der Worterkennung beim kompetenten Leser.....	17
2.2.1 Dual-Route-Cascaded Model.....	18
2.2.2 Triangel-Modell.....	21
2.3 Morpheme als Einheiten der deutschen Schriftsprache.....	24
2.3.1 Definition von Morphemen.....	24
2.3.2 Prinzip der Morphemkonstanz.....	26
2.4 Theoretische Modelle zur Verarbeitung von Morphemen in der Worterkennung. .	28
2.5 Entwicklung des Wissens um morphologische Einheiten.....	34
2.6 Verarbeitung von Morphemen beim lauten Lesen.....	38
3 Empirie.....	43
3.1 Fragestellung und Hypothesen.....	43
3.2 Methode.....	44
3.2.1 Beschreibung der TeilnehmerInnen.....	44
3.2.2 Design der Untersuchung.....	45
3.2.2.1 Leseaufgabe: Pseudowörter.....	45
3.2.2.2 Leseaufgabe: Störung.....	46
3.2.3 Durchführung.....	47

3.2.4 Analyse und Kodierung der Tonspuren.....	49
3.3 Ergebnisse.....	50
3.3.1 Ergebnisse der Pseudowortbedingungen.....	51
3.3.2 Ergebnisse der Trennbedingungen.....	55
4 Diskussion.....	58
4.1 Methode.....	58
4.2 Ergebnisse.....	61
4.3 Ausblick.....	67
5 Zusammenfassung.....	69
5.1 Deutsche Zusammenfassung.....	69
5.2 English Summary.....	70
Literaturverzeichnis.....	73
Anhang A: Pseudowort-Listen.....	76
Anhang B: Wortlisten der Trennbedingung.....	77
Anhang C: Lebenslauf.....	79

Abbildungsverzeichnis

Abbildung 1: Fixierungs-/Orientierungspunkt am Bildschirm.....	47
Abbildung 2: Präsentation des Stimulus am Bildschirm.....	48
Abbildung 3: Profildigramm der geschätzten Randmittelwerte für die Gesamtlesezeiten je Pseudowortbedingung (1=Pseudowort aus Morphemen, 2=abgeleitetes Pseudowort ohne Morpheme) und Schulklasse.....	54
Abbildung 4: graphische Darstellung der mittleren Lesezeiten je Trennbedingung in Sekunden.....	57

Tabellenverzeichnis

Tabelle 1: Mittlere Lesegesamtzeiten (in Sekunden, auf Tausendstel gerundet) nach Geschlecht, Standardabweichung, Anzahl der Items (unter Ausschluss der ungültigen Lesezeiten).....	50
Tabelle 2: Mittlere Lesegesamtzeiten (in Sekunden, auf Tausendstel gerundet) nach besuchter Klasse, Standardabweichung, Anzahl der Items (unter Ausschluss der ungültigen Lesezeiten).....	51
Tabelle 3: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Pseudowortbedingung, Standardabweichung und Anzahl der SchülerInnen....	52
Tabelle 4: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Pseudowortbedingung nach Geschlecht geteilt, Standardabweichung und Anzahl der SchülerInnen.....	53
Tabelle 5: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Pseudowortbedingung nach Klasse, Standardabweichung und Anzahl der SchülerInnen.....	54
Tabelle 6: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Trennbedingung, Standardabweichung und Anzahl der SchülerInnen.....	55

1 Einleitung

Die Ergebnisse der Pisa-Studien sorgen wiederholt für großes Medieninteresse und Entsetzen über die Lesekompetenz von österreichischen SchülerInnen. Dies führt zu Bemühungen den Unterricht an Schulen zu verbessern und zu angeregter Forschung. Nicht nur das Verständnis um das Gelesene, sondern auch grundlegende Lesefertigkeiten tragen zu kompetentem Lesen bei. So weisen schwache, auch ältere Leser nicht nur eine langsame Verarbeitung beim Lesen auf, sondern auch eine mangelnde Fähigkeit Wörter zu dekodieren (Noack, 2006; Schmidt-Barkow, 2002; zitiert nach Bredel et al., 2011). Zu den Einheiten, welche genutzt werden können, um Wörter zu dekodieren, zählen neben Buchstaben auch Silben und Morpheme. Morphologische Einheiten werden gegen Ende der Grundschulzeit zwar vermittelt, jedoch lassen sich empirisch fundierte Annahmen, welche ein Vermitteln dieser Einheiten als für den Leseprozess förderlich begründen, im deutschen Sprachraum erst vereinzelt finden (beispielsweise Landerl & Reitsma, 2005).

Studien im italienischen Sprachraum zeigen, dass Kinder bereits in der 3. Klasse Grundschule (Burani et al. 2002) morphologische Einheiten im Prozess des lauten Lesens nutzen und dies die Lesegeschwindigkeit erhöht. Jedoch können diese Ergebnisse trotz der gemeinsamen orthographischen Transparenz nicht unmittelbar auf den deutschen Sprachraum übertragen werden. Im deutschen Sprachraum gibt es nach Kenntnisstand des Autors der vorliegenden Arbeit bis dato keine empirischen Befunde aus Studien, welche die Verarbeitung von morphologischen Strukturen beim lauten Lesen an Kindern untersuchen.¹

Die vorliegende Arbeit geht der Frage nach, welche Verwendung morphologische Strukturen im Prozess des lauten Lesens bei Kindern am Ende der vierten Klasse Grundschule finden. Dadurch soll ein Beitrag zu grundsätzlichen Fragestellungen betreffend der Nutzung von Einheiten in einem erfolgreichen Leseprozess geleistet werden und weitere Forschungen angeregt werden, mit dem Ziel anhand dieser Forschungen Implikationen für die schulische und lehrende Praktik im weitesten Sinn abzuleiten, um im Besonderen Leseanfänger und schwache Leser zu unterstützen.

¹ Abgesehen von einer parallel zu dieser erfolgten Studie von Barbara Sonnleitner an Kindern der 2. Klasse Grundschule, welche aufgrund einer gemeinsamen Entwicklung des Designs dieselbe Vorgehensweise und Methode wie die vorliegende Arbeit aufweist.

2 Theoretische Grundlagen

Die theoretischen Grundlagen der vorliegenden Arbeit umfassen im ersten Schritt allgemeine Theorien der Leseentwicklung und des Worterkennens. Danach wird genauer auf die Rolle des Morphems in der deutschen Sprache, auf Modelle der morphologischen Verarbeitung, die Entwicklung des Wissens um morphologische Strukturen und Studien zur morphologischen Verarbeitung beim lauten Lesen eingegangen. In den folgenden Modellen und Theorien wird wiederholt auf die Bedeutung von Unterschieden in den jeweiligen Schriftsprachsystemen für die Leseentwicklung sowie die Worterkennung hingewiesen. Um dem Leser der vorliegenden Arbeit ein besseres Verständnis zu ermöglichen, werden im Folgenden die Struktur der deutschen Schriftsprache und deren Begrifflichkeiten kurz dargestellt.

Die deutsche Schriftsprache basiert auf einem orthographisch-transparenten Schriftsystem und weist eine hohe Regularität von Orthographie-zu-Phonologie Korrespondenzen auf (Coltheart et al., 2000). Das bedeutet, dass die Verbindung zwischen gesprochener und geschriebener Sprache nicht durch die einfache Zuordnung von Buchstabe zu Laut, jedoch auf der abstrakten Ebene von Phonem (die kleinste bedeutungsdifferenzierende Einheit des Lautsystems basierend auf der Zusammenfassung von Lauten zu Klassen) und Graphem (die kleinste Einheit des Schriftsystems basierend auf Buchstaben oder Buchstabengruppen, die mit einem Phonem korrespondieren) relativ konsistent abgebildet werden kann (Schröder-Lenzen, 2013, S. 15ff). Auch wenn die Phonem-Graphem-Korrespondenz nicht in allen Fällen eindeutig ist, kann die Verschriftlichung der deutschen Sprache als lautgetreu bezeichnet werden. 73% der Laute werden durch die häufigsten Buchstaben repräsentiert (Naumann, 1989; zitiert nach Schröder-Lenzen, 2013, S. 27).

2.1 Entwicklungsmodelle des Schriftspracherwerbs

„Es gibt keine Stunde Null für den Schriftspracherwerb – weder in dem Sinne, dass alle Kinder die gleichen Voraussetzungen hätten, noch in dem Sinne, dass alle Kinder gar nichts wüssten.“ (Günther, 2010, S. 16). Diese Einschätzung findet sich in Entwicklungsmodellen, wie etwa dem Kompetenzentwicklungsmodell nach Klicpera,

Schabmann und Gasteiger-Klicpera (2010, S. 25ff). Unterschiede in vorschulischen Kompetenzen, wie beispielsweise die Phonologische Bewusstheit², werden als bedeutsame, jedoch nicht alleinige Komponenten für die spätere Leseentwicklung gesehen (Klicpera et al., 2010, S.28).

Im Folgenden werden Theorien der Leseentwicklung beschrieben, deren Anwendbarkeit auf die deutsche Sprache gerechtfertigt scheint. Diese berücksichtigen empirische Befunde im deutschen Sprachraum (Kompetenzentwicklungsmodell, Klicpera et al., 2010, S. 25ff) beziehungsweise können entsprechend moduliert werden (Modell des Sichtwortlesens, Ehri, 2005) oder beachten bereits in ihrer Entwicklung sprachspezifische Unterschiede und sollen diese erklären (linguistische Grain Size Theorie, Ziegler & Goswami, 2005).

2.1.1 Phasenmodell des Sichtwortlesens

Das Phasenmodell von Ehri (2005; Ehri & McCormick, 1998) soll Prozesse erklären, welche Kinder in der Entwicklung der Lesefähigkeit durchlaufen, mit der Absicht, dass durch die Kenntnis dieser Entwicklungsphasen insbesondere die Schwierigkeiten leseschwacher Kinder verstanden und durch PädagogInnen unterstützt werden können. Ehri (2005) unterscheidet allgemein zwischen Lesestrategien, die von geübten Lesern angewandt werden können, um unbekannte Wörter zu lesen (z.B.: Phonologisches Rekodieren) und solchen, die Anwendung beim Lesen bekannter Wörter finden. Die als Sichtwortlesen bezeichnete Lesestrategie beschreibt das automatisierte Lesen bekannter Wörter mittels direkten Zugangs zu einer im Gedächtnis gespeicherten Repräsentation dieses Wortes. Diese Repräsentation enthält Information zu Aussprache sowie Bedeutung des Wortes. Dieser Weg stellt nach Ehri und McCormick (1998) den schnellsten, effizientesten und am wenigsten fehleranfälligen Prozess dar und erlaubt dem Leser den Fokus auf der Bedeutung des Gelesenen zu halten. Jedes Wort, welches ausreichend oft gelesen wurde, kann schließlich aus dem Gedächtnis abgerufen werden und zunehmend automatisch, dass heißt ohne bewusste Aufmerksamkeit, gelesen werden. Somit kommt

² Phonologische Bewusstheit meint die Fähigkeit, jegliche phonologische Einheit (z.B.: Silben, einzelne Phoneme) innerhalb eines Wortes zu erkennen, zu identifizieren und zu manipulieren (Ziegler & Goswami, 2005). Phonologische Bewusstheit und Erwerb der Schriftsprache kann als ein interaktiver Prozess gesehen werden. Durch Verständnis von Elementen der schriftlichen Sprache (Buchstaben, Wörter) können Elemente der mündlichen Sprache erst entdeckt und strukturiert werden, dies führt wiederum zu einem besseren Verständnis der Schrift (Günther, 2010, S. 16).

dem Aufbau eines Sichtwortschatzes große Bedeutung zu. (Ehri, 2005; Ehri & McCormick, 1998)

Anhand von vier (Ehri, 2005) beziehungsweise fünf (Ehri & McCormick, 1998) Phasen der Leseentwicklung wird die jeweils vorherrschend angewandte Strategie des Lesens beschrieben. Das Anwenden einer Strategie bedeutet nicht das Ausschließen einer anderen, vielmehr handelt es sich um eine kontinuierliche Entwicklung, die auch durch das Vorhandensein unterschiedlicher Lesestrategien innerhalb einer Phase charakterisiert ist. Diese Phasen werden wie folgt charakterisiert:

Die *präalphabetische Phase* beschreibt keinen Leseprozess im eigentlichen Sinn, da Kinder in diesem frühen Stadium noch kaum etwas über alphabetische Prinzipien und über Verbindungen zwischen Buchstaben und Lauten wissen. Wörter werden aufgrund von visuellen Merkmalen beziehungsweise dem Kontext erkannt (z.B. Logos bekannter Marken). Die präalphabetische Phase charakterisiert Kinder vor Schuleintritt beziehungsweise stark beeinträchtigte Leser.

In der *partiell alphabetischen Phase* beginnen Kinder Zusammenhänge zwischen einigen Buchstaben und deren Aussprache zu lernen und diese Fähigkeit zu nutzen, um Wörter zu erlesen. Jedoch fehlen vollständige Kenntnisse über Graphem-Phonem-Verbindungen. Die Fähigkeit Wörter zu dekodieren ist somit nur rudimentär ausgeprägt. Diese Phase beschreibt Kinder zu Beginn des Schuleintritts beziehungsweise schwache Leser.

Die *voll alphabetische Phase* beschreibt das Erlangen vollständiger alphabetischer Kenntnisse und des umfassenden Wissens über Graphem-Phonem-Zuordnungen. Kinder können auch unbekannte Wörter lesen. Wörter werden zunehmend genauer gelesen, von ähnlichen Wörtern unterschieden, wieder erkannt und deren Verbindungen im Gedächtnis gestärkt. Ab der ersten Klasse sollten SchülerInnen üblicherweise in die voll und ab der zweiten in die konsolidierte alphabetische Phase übergehen.

In der *konsolidierten alphabetischen Phase* wird das Gedächtnis der Kinder kontinuierlich um Sichtwörter erweitert. Graphem-Phonem-Verbindungen werden zu größeren Einheiten (z.B. Silben, Morphemen, ganzen Wörtern) zusammengefügt, verdichtet und im Gedächtnis gespeichert. Aufgrund dieser größeren

Speichereinheiten genügt nun das Aktivieren einer verhältnismäßig geringen Anzahl von Verbindungen, um auch mehrsilbige Wörter zu lesen. Der Leseprozess wird zunehmend automatisiert.

Die *automatisierte alphabetische Phase* (Ehri & McCormick, 1998) beschreibt den reifen Leser, der die meisten Wörter automatisiert mittels direkten lexikalischen Zugangs beziehungsweise unbekannte Wörter mithilfe der effizientesten Strategie erfassen kann.

Das beschriebene Phasenmodell wurde im englischsprachigen Raum entwickelt und es stellt sich die Frage der Übertragbarkeit auf andere Sprachen, insbesondere jene mit transparenten Orthographien (Ehri, 2005). Nach Klicpera et al. (2010, S.27) ist die vorherrschende Strategie der präalphabetischen Phase bei Kindern im deutschen Sprachraum nicht zwingend vorhanden. Dennoch scheint das Modell die Entwicklung der Lesefertigkeiten im deutschen Sprachraum abbilden zu können. Dies bezieht sich besonders auf die parallele Entwicklung von Lesefertigkeiten auf Wortebene und Buchstabenebene (Klicpera & Gasteiger-Klicpera, 1999).

2.1.2 Kompetenzentwicklungsmodell des Lesens

Das Kompetenzentwicklungsmodell des Lesens (Klicpera, Schabmann & Gasteiger-Klicpera, 2010, S. 25ff) grenzt sich von gängigen Entwicklungsmodellen ab, indem es vor allem den Erwerb von notwendigen Lesekompetenzen und weniger das Durchlaufen bestimmter Phasen oder Stufen als maßgeblich für die Leseentwicklung ansieht. Die Entwicklung der Lesefähigkeit wird in Abhängigkeit von individuellen Leistungsbeziehungsweise Lernvoraussetzungen der Kinder und der im Unterricht beziehungsweise im schulischen Umfeld erfolgten Leseinstruktion gesehen. Vorschulische Kompetenzen umfassen phonologische Bewusstheit, Gedächtnisleistung und Steuerung der visuellen Aufmerksamkeit. Diese können die spätere Leseentwicklung, wenn auch nur teilweise und unter bestimmten Bedingungen (an Erstleseinstruktion geknüpft), vorhersagen. Im Sinne der *präalphabetischen Phase* nach Ehri (1998; 2005) werden auch Versuche von einigen Kindern genannt, Wörter anhand visueller Merkmale zu erraten.

Der Eintritt in die Schule, der mit ersten Leseinstruktionen verbunden ist, wird laut Klicpera et al. (2010) als kritischer Moment der Leseentwicklung betrachtet. Vor

Schuleintritt verfügen Kinder im deutschsprachigen Raum meist nur über geringe alphabetische Kenntnisse. Der Schulanfang initiiert den Beginn des „echten“ Lesenlernens und den Übergang zur *alphabetischen Phase mit geringer Integration*. Kinder erlernen das alphabetische Prinzip, können Buchstaben und Laute einander zuordnen und erlangen die Fertigkeit der phonologischen Rekodierung, um Wörter zu lesen. Da dies aufgrund der regelmäßigen Phonem-Graphem-Verbindungen in der deutschen Sprache eine zuverlässige Strategie ist, entwickeln Kinder bereits nach wenigen Wochen erste Lesekompetenzen. In Abhängigkeit von der Lautorientierung des Leseunterrichts können Kinder ihre Fertigkeit zum phonologischen Rekodieren mehr oder weniger rasch ausbauen und diese begünstigt den Aufbau eines mentalen Lexikons. Die lexikalische und die nicht-lexikalische Lesestrategie (siehe auch den Abschnitt zu Zwei-Wege-Modelle in dieser Arbeit, S.18ff) entwickeln sich nicht unabhängig voneinander und stehen in starker Interaktion mit der Leseinstruktion. Mit zunehmender Automatisierung des Lesevorgangs und abhängig von individuellen Gedächtniskomponenten erhöht sich die Fähigkeit Wörter direkt aus dem mentalen Lexikon abzurufen. (Klicpera et al., 2010, S. 25ff)

In der *alphabetischen Phase mit voller Integration* wird sowohl die Fähigkeit zum phonologischen Rekodieren als auch die des direkten Zugangs zum mentalen Lexikon automatisiert. Dies führt zu einer Erhöhung der Lesegeschwindigkeit und sinkenden Anzahl von Lesefehlern. Zusätzlich scheint sich der Prozess des phonologischen Rekodierens zu beschleunigen, indem Kinder ein neues, unbekanntes Wort nicht mehr Buchstabe für Buchstabe erlesen, sondern größere mentale Einheiten, etwa häufig auftretende Buchstabenabfolgen, zur Verarbeitung nutzen. Die beschriebenen Lesestrategien interagieren immer stärker miteinander und die Entscheidung, mit welcher Strategie ein Wort am effizientesten gelesen werden kann, läuft automatisch ab. Dies kennzeichnet den Übergang zur letzten Phase der Leseentwicklung, der *automatisierten und konsolidierten Integration aller beteiligten Verarbeitungsprozesse*. (Klicpera et al., 2010, S. 25ff)

Die Modelle der Leseentwicklung von Klicpera et al. (2010) und Ehri (2005) stimmen in wesentlichen Punkten überein. Jedoch wird im Kompetenzentwicklungsmodell (Klicpera et al., 2010) gezielt auf den Leseerwerb im deutschen Sprachraum eingegangen und der Einfluss der (schulischen) Instruktion auf die Leseentwicklung miteinbezogen. Des Weiteren steht nicht der (zeitliche) Ablauf von Phasen im Vordergrund, sondern das

Erlangen von Kompetenzen. Leseentwicklung bewegt sich also zur Kompetenz hin, je nach gestellter Anforderung auf die dafür effizienteste Strategie zurückgreifen zu können, um diese erfolgreich zu bewältigen.

2.1.3 Psycholinguistische Grain Size Theorie

Die Grain Size Theorie (Ziegler & Goswami, 2005) stellt ein theoretisches Rahmenwerk des Lesens und seiner Entwicklung zur Verfügung, welches die empirisch beobachteten sprachspezifischen Unterschiede im Leseerwerb erklären soll. Die Entwicklung der Lesefertigkeit unterscheidet sich in den jeweiligen Sprachen aufgrund der Konsistenz der Graphem-Phonem-Beziehung, der Größe der orthografischen und phonologischen Verarbeitungseinheiten und aufgrund der angewandten Lehrmethoden. „Grain Sizes“ bezeichnen diejenigen Einheiten von lexikalischen Repräsentationen, welche das jeweilige Schriftsystem am beständigsten abbilden können und dem kompetenten Leser zur Verfügung stehen. Zu Beginn des Leseerwerbs müssen Kinder diese Einheiten jedoch erst entdecken. (Ziegler & Goswami, 2005)

Nach Ziegler und Goswami (2005) werden Leseanfänger mit folgenden Problemen konfrontiert, deren Lösung sich je nach der zu lernenden Schriftsprache unterscheidet und den Leseerwerb beeinflusst:

Das *Verfügbarkeitsproblem* beschreibt, dass nicht automatisch alle phonologischen Einheiten bereits vor dem Schriftspracherwerb verfügbar sind, sondern erst entwickelt werden müssen.

Das *Konsistenzproblem* ergibt sich, da einige orthographische Einheiten auf unterschiedliche Art und Weise ausgesprochen werden können beziehungsweise umgekehrt manchen phonologischen Einheiten mehrere Schreibweisen entsprechen. Dies variiert sowohl nach Sprache als auch nach den jeweiligen orthographischen Einheiten und wirkt sich auf die Entwicklung des Lesens aus. Die Fähigkeit zum phonologischen Rekodieren bildet sich in Orthographien, die durch eine hohe Konsistenz von Buchstaben-Laut-Beziehungen gekennzeichnet sind, früher aus als in inkonsistenten Orthographien.

Das *Granularitätsproblem* bezeichnet den Umstand, dass mehrere orthographische Einheiten gelernt werden müssen, wenn größere Einheiten notwendig sind, um phonologische Repräsentationen abzubilden. Dies bedeutet vereinfacht, dass es weniger Buchstaben zu lernen gilt als deren mögliche Kombinationen. Von der Struktur der jeweiligen Schriftsprache ist abhängig, welche Einheiten von lexikalischen Repräsentationen Kinder lernen müssen, um effizient, genau und fehlerfrei lesen zu können. In orthographisch konsistenten Sprachen (z.B. Deutsch, Griechisch, Italienisch, Spanisch) können sich Leseanfänger eher auf kleinere Einheiten verlassen, um Wörter richtig zu erkennen und auszusprechen als in inkonsistenten Orthographien (z.B.: Englisch, Französisch). Das Entwickeln von größeren Einheiten erfordert das Lernen von weitaus mehreren orthographischen Einheiten und ist somit mit größerer Anstrengung und einem langsameren Fortschritt im Leseerwerb verbunden.

Kulturübergreifende empirische Studien zeigen, dass Schüler Ende der ersten Klasse Grundschule in konsistenten Orthographien beim Lesen von Wörtern und Pseudowörtern eine weitaus höhere Genauigkeit aufweisen als in inkonsistenten Orthographien (Seymour et al., 2003; zitiert nach Ziegler & Goswami, 2005). In konsistenten Orthographien scheint das Bilden größerer und folglich mehrerer Einheiten nicht notwendig, da sie länger und aufwendiger scheinen als einfache Graphem-Phonem-Zuordnungen. Im Sinne der Grain-Size-Theorie ist somit in der deutschen Sprache das Entwickeln von „Grain Sizes“ in der Größe von Morphemen im Laufe des Schriftspracherwerbs nicht unbedingt erforderlich, um Wörter fehlerfrei lesen und aussprechen zu können. Goswami und Ziegler (2006) ergänzen jedoch, dass dem kompetenten Leser auch größere Einheiten von lexikalischen Repräsentationen zu Verfügung zu stehen, um den Leseprozess zu beschleunigen und zu erleichtern. Große Einheiten, die Verbindungen zwischen Orthographie, Phonologie und Semantik herstellen, garantieren insbesondere einen schnellen Zugang zur Bedeutung des Gelesenen. Der Einfluss der Morphologie einer Sprache auf die phonologische Entwicklung wurde zwar nicht explizit in die Theorie eingebaut, diese kann jedoch sowohl in Bezug auf die Verfügbarkeit von Einheiten lexikalischer Repräsentationen vor dem Leseerwerb als auch in Dekodierprozessen eine Rolle spielen (Goswami & Ziegler, 2006).

In den Entwicklungstheorien von Ehri (2005) und Klicpera et al. (2010) wird das Zusammenfassen von Graphem-Phonem-Verbindungen zu größeren Einheiten im Laufe der Leseentwicklung explizit erwähnt. Dadurch kann sich in der konsolidierten

alphabetischen Phase (Ehri, 2005) beziehungsweise der alphabetischen Phase mit voller Integration (Klicpera et al., 2010) die Lesegeschwindigkeit erhöhen, insbesondere beim phonologischen Rekodieren. SchülerInnen der vierten Klasse Grundschule, welche die Stichprobe der vorliegenden Arbeit bilden, sollten laut diesen Modellen häufige Buchstabenmuster als Einheiten rekodieren können. Zu diesen Einheiten werden neben Wörtern und Silben auch Wortstämme und Suffixe, somit Morpheme gezählt (Ehri, 2005; Ehri & McCormick, 1998).

2.2 Modelle der Worterkennung beim kompetenten Leser

Modelle der Worterkennung beschäftigen sich mit der Frage, wie ein Wort von kompetenten Lesern, also nach erfolgreicher Schriftsprachentwicklung, erkannt, dekodiert und im Falle des lauten Lesens artikuliert wird. Die im Folgenden angeführten Modelle vertreten die gemeinsame Grundannahme, dass spezielle kognitive Einheiten für die Verarbeitung von Schrift zuständig sind, welche über jeweilige Verbindungen einen erfolgreichen Leseprozess ermöglichen. Weitere Komponenten, die den Prozess des Wortlesens beeinflussen können (Umgebungsfaktoren, Instruktion im Unterricht, Förderung im außerschulischen Bereich, sozialer Kontext, etc.) finden in kognitiv orientierten Modellen keine Beachtung, spielen jedoch bezüglich der Entwicklung des Leseerwerbs bei Kindern eine bedeutende Rolle (Klicpera et al., 2010, S.42f). Die Modelle beschränken sich auf die Ebene des Wortes, geben also keine Auskunft darüber, wie Sätze beziehungsweise ganze Texte verarbeitet werden. Trotz dieser Einschränkung kommt diesen Modellen große Bedeutung zu, da Wortlesestrategien eine zentrale Rolle im Leseprozess einnehmen. Sie sind für die Erklärung erfolgreicher Verarbeitung der Schriftsprache wichtig und erlauben dadurch Rückschlüsse, welche Verarbeitungskomponenten bei einer Lese- Rechtschreibschwäche beeinträchtigt sein können. (Klicpera et al., 2010, S. 42f; Steinbrink & Lachmann, S.29)

Grundsätzlich dominieren zwei verschiedene Ansätze den psycholinguistischen Forschungsbereich: die Zwei-Wege-Modelle und die Netzwerkmodelle. In Zwei-Wege-Modellen werden zwei getrennte Wege der Wortverarbeitung angenommen, die direkte Worterkennung bekannter Wörter und die serielle Rekodierung unbekannter oder Pseudo-Wörter mittels der Regeln der Graphem-Phonem-Zuordnung (Klicpera et al. 2010, S. 44f;

Steinbrink & Lachmann, 2014, S. 29f). Die beiden Wege wurden in früheren Zwei-Wege-Modellen als einander ausschließend betrachtet, eine Annahme, die nach heutigem Forschungsstand nicht mehr aufrechterhalten wird (Klicpera et al., 2010, S. 44). In Netzwerkmodellen wird nur ein Verarbeitungsweg angenommen. Man kann somit von einer Theorie des einfachen Zugangsweges sprechen (Klicpera et al., 2010, S. 48). Des Weiteren erfolgt bei der Worterkennung kein Zugriff auf ein mentales Lexikon im Sinne der Zwei-Wege-Theorien. Auch das Anwenden von expliziten Regeln der Graphem-Phonem-Korrespondenz wird in Frage gestellt (Klicpera et al., 2010, S.48; Plaut et al., 1996). Ein Wortstimulus führt zur Aktivierung von Mustern, welche die Verbindungen der jeweiligen Buchstabenabfolge im Netzwerk repräsentieren (Plaut et al., 1996).

Beide Modelle werden mittels Computersimulationen hinsichtlich der Übertragbarkeit der jeweiligen Annahmen auf menschliche Leseprozesse überprüft. Beide Modelle können zahlreiche empirisch beobachtbare Effekte des Wortlesens erklären, jedoch bleiben durch die Beschränkung auf einsilbige Wörter Fragen bezüglich des Lesens komplexer, mehrsilbiger Wörter offen. Aus diesem Grund erfahren die jeweiligen Modelle laufend Erweiterungen und nähern sich zunehmend einander an (Steinbrink & Lachmann, 2014, S. 43).

Im Folgenden werden das Dual-Route-Cascaded-Model nach Coltheart et al. (2001) als Vertreter der Zwei-Wege-Modelle und das Triangel-Modell nach Plaut et al. (1996) als Vertreter der Netzwerkmodelle ausführlicher beschrieben.

2.2.1 Dual-Route-Cascaded Model

Das Dual-Route-Cascaded-Model (DRC-Modell; Coltheart, Rastle, Perry, Langdon & Ziegler, 2001) stellt eine Weiterentwicklung der Zwei-Wege-Modelle (Coltheart, 1978; McClelland & Rumelhart, 1982; zitiert nach Coltheart et al., 2001) dar. Es handelt sich um ein computerisiertes Modell der visuellen Worterkennung und des lauten Lesens, welches auch in seiner Erweiterung auf den deutschen Sprachraum (Coltheart et al., 2000) die menschliche Leseleistung zufriedenstellend abbilden konnte. Jedoch gilt dies wie im englischen Modell nur für einsilbige Wörter, die sich aus maximal acht Buchstaben zusammensetzen.

Im Gegensatz zu früheren Zwei-Wege-Modellen wird die Annahme vertreten, dass ein Wort beziehungsweise Pseudowortstimulus zu einer Aktivierung beider Routen führen kann. Insbesondere beim Lesen wortähnlicher Pseudowörter wird ersichtlich, dass zwei einander ausschließende Routen den Leseprozess nicht ausreichend abbilden können (Pseudowörter, die eine wortähnliche Graphemfolge zeigen, werden schneller gelesen als solche, die vollkommen wortunähnlich sind). Dies leitet zu einer weiteren Grundannahme des Modells über, der unmittelbaren, stufenweisen („cascaded“) Aktivierung: Die geringste Aktivierung einer Verarbeitungsstufe wird unmittelbar an die darauf folgende weitergeleitet, ohne einen gewissen Schwellenwert erreichen zu müssen. (Coltheart et al., 2001)

Jede Route setzt sich aus miteinander interagierenden Stufen zusammen, welche wiederum ein Set von Einheiten („units“) beinhalten. Einheiten repräsentieren die kleinsten symbolischen Teile des DRC-Modells, beispielsweise Repräsentationen von Wörtern im orthographischen Lexikon. Die Einheiten angrenzender Stufen sowie innerhalb derselben Ebene interagieren über gegenseitige Aktivierung oder Hemmung kompatibler beziehungsweise inkompatibler Einheiten, bis auf zwei Ausnahmen: Die Verbindung von Einheiten des orthographischen Lexikons zu jenen des phonologischen Lexikons erfolgt nur über Aktivierung und die Verbindung von der Ebene der visuellen Merkmale geht nur in Richtung der Buchstabenebene. (Coltheart et al., 2001)

Der Prozess des lauten Lesens beginnt nach Coltheart et al. (2001) immer mit der Aktivierung einer Buchstabenabfolge auf der Ebene der visuellen Merkmale und endet mit der Aktivierung dieser Buchstabenabfolge auf der Phonem-Ebene. Die visuellen Merkmale der Buchstaben aktivieren abstrakte Buchstabenrepräsentationen, welche bereits von bestimmten visuellen Merkmalen (Groß-/Kleinschreibung, Format, etc.) gelöst sind, jedoch die Position der Buchstaben zueinander beibehalten. Danach teilt sich der Weg laut Coltheart et al. (2001) in die lexikalische und die nicht-lexikalische Route:

Über die lexikalischen Route wird der Eintrag des Wortes im orthographischen Lexikon aktiviert. Dieser aktiviert den zugehörigen Eintrag im phonologischen Lexikon, welcher die Phoneme des Wortes (sowie deren Position zueinander) aktiviert. Das orthographische Lexikon enthält die schriftliche Form aller Wörter, deren Schreibweise dem Leser bekannt ist. Das phonologische Lexikon ist nach demselben Prinzip aufgebaut und enthält die gesprochene Form. In beiden Lexika existiert je ein

Eintrag pro Wort, welche jeweils äquivalent zueinander sind. Die Einheiten des orthographischen Lexikons sind empfänglich für Häufigkeiten. Hochfrequente Wörter führen zu einer schnelleren Aktivierung der jeweiligen Repräsentation des Wortes als niederfrequente Wörter. Über die lexikalische Route ist auch der Zugriff auf das semantische Lexikon möglich, welches die Bedeutung des Wortes beinhaltet und sowohl mit dem orthographischen als auch dem phonologischen Lexikon in bidirektionaler Verbindung steht.

Über die nicht-lexikalische Route wird eine Buchstabenfolge anhand der Regeln der Graphem-Phonem-Korrespondenz (GPK) in eine Phonemfolge umgewandelt. Diese aktiviert nun die einzelnen entsprechenden Phoneme der Phonemebene und die Aussprache folgt. Die Konvertierung der Buchstabenfolge erfolgt sukzessive von links nach rechts durch das Finden einer passenden Regel. Eine Grundregel besagt, dass immer ein Graphem, sei es ein Buchstabe oder ein Buchstabencluster, einem Phonem entspricht. Mehrzeichenregeln („multiletter-rules“) legen fest, welche Buchstabenkombinationen zwischenzeitlich zu Einheiten zusammengefasst werden, und sind positions- und kontextspezifisch. Laut Coltheart, Perry und Ziegler (2000) werden für den deutschen Sprachraum aufgrund der orthographischen Transparenz weniger und kleinere Mehrzeichenregeln als im Englischen benötigt, zudem zeigen sich diese verlässlicher. Jedoch werden mehrere kontextspezifische Regeln benötigt, die eine korrekte Aussprache eines Graphems nur unter Beachtung vorheriger oder folgender Buchstaben erlauben.

Die Verwendung von gespeicherten morphologischen Einheiten im Leseprozess scheint nach dem Dual-Route-Cascaded-Model (Coltheart et al., 2001) über beide Routen eher unwahrscheinlich. Entlang der lexikalischen Route wird ein Zugriff auf Repräsentationen des gesamten Wortes und entlang der nicht-lexikalischen Route nur das zwischenzeitliche Zusammenfassen von Buchstaben zu Einheiten angenommen. Jedoch bleibt diese Annahme auf Leseprozesse von einsilbigen Wörtern beschränkt. Rastle und Coltheart (2000) verfolgten in einer explorativen Studie, wie ein Regelsystem für zweisilbige Wörter im Sinne der Zwei-Wege-Theorie aussehen müsste, welches gewährleisten kann, dass Wörter als auch Pseudowörter richtig betont ausgesprochen werden. Mehrsilbige Wörter erfordern eine Beachtung von Silbenanlaut und -auslaut. Dies kann im nicht-lexikalischen System nur über den Zugriff auf gespeicherte Einheiten gewährleistet werden, die

Buchstaben zu Suffixen, somit morphologischen Einheiten kombinieren. Auch wenn gezeigt werden konnte, dass der Leseprozess von zweisilbigen Wörtern mit einem Regelsystem ausgedrückt werden kann, ist dieses für eine Implementierung in das DRC-Modell noch zu wenig ausgereift. (Rastle & Coltheart, 2000)

2.2.2 Triangel-Modell

Das Triangel-Modell nach Plaut, McClelland, Seidenberg und Patterson (1996) stellt ein netzwerkbasierendes Modell der Worterkennung dar. Im Gegensatz zu Zwei-Wege-Modellen werden in diesen konnektionistischen Modellen keine getrennten Mechanismen für die Bearbeitung von bekannten und unbekannten Wörtern beziehungsweise Pseudowörtern angenommen. Jede Buchstabenfolge wird mittels des gesamten Systems bearbeitet. Die Verarbeitung einer Buchstabenabfolge verläuft beim kompetenten Leser parallel. Das heißt, es wird nicht jeder Buchstabe seriell und einzeln verarbeitet, sondern immer mehrere Buchstaben parallel. Auch auf Phonem und auf Graphem-Ebene findet die Verarbeitung parallel statt. (Plaut et al., 1996)

Dem Netzwerksystem liegt nach Plaut et al. (1996) ein Lernprozess zugrunde. Im Leseprozess werden Verbindungen anhand der Interaktion bestimmter Aktivitätsmuster über Gruppen von Einheiten hinweg gelernt. Wörter werden somit nicht als Einheit in einem mentalen Lexikon gespeichert, sondern in ihren erlernten Verbindungen über interagierende Einheiten, welche wiederum an Leseprozessen von vielen anderen Wörtern beteiligt sind (im Sinne der Netzwerk-Modelle bezeichnet dies die Annahme der verzweigten Repräsentationen). Diese Einheiten können sich gegenseitig in ihrem Aktivierungslevel beeinflussen und stellen kodierte orthographische, phonologische und semantische Informationen zur Verfügung, welche über sogenannte „hidden units“ vermittelt werden. Diese verborgenen Einheiten kodieren über Gewichtung der Verbindungen das Wissen des Systems, wie die unterschiedlichen Typen von Information miteinander in Beziehung stehen. Die spezifischen Werte der Gewichtung erfolgen aufgrund eines automatischen, nicht instruierten Lernprozesses, der allein durch die Konfrontation des Systems mit geschriebenen Wörtern und deren Aussprache initiiert wird. (Plaut et al., 1996)

Die Leistung des Netzwerks entwickelt und verbessert sich zunehmend über Anpassung und Lernalgorithmen. Das heißt, das System versucht zu einem „korrekten Ergebnis“ zu gelangen (im Falle des lauten Lesens zu einer fehlerfreien Aussprache), indem die Verbindungen zwischen den Einheiten mittels Rückübertragung angepasst und deren Gewichtungen verändert werden. Des Weiteren ist das Netzwerk empfänglich für die Häufigkeit des Auftretens eines Wortes. Wiederholte Aktivierung von Verbindungen führt dazu, dass diese Aktivitätsmuster gestärkt werden. Bei zukünftigen Aktivierungen können diese Verbindungen schneller und eher fehlerfrei abgerufen werden. Da ähnliche Wörter durch ähnliche Aktivitätsmuster repräsentiert werden (Annahme der verzweigten Repräsentationen), entwickeln sich diese gemeinsam. Bestimmte Verbindungen von Buchstaben, die häufig erfahren werden, werden als Muster gelernt und gespeichert. Diese können schließlich in zahlreichen ähnlichen Wörtern, die dieses Muster aufweisen, schneller abgerufen werden und somit den Leseprozess beschleunigen. Sowohl die Konsistenz (Aussprache eines Wortes stimmt mit ähnlich geschriebenen Wörtern überein) als auch die Worthäufigkeit beeinflussen folglich die Dauer der Worterkennung. (Plaut et al., 1996)

Das Triangel-Modell (Plaut et al., 1996) setzt sich aus drei Ebenen zusammen, wobei die Ebene der orthographischen Einheiten über die Ebene der „hidden units“ mit der Ebene der phonologischen Einheiten verbunden wird. Im Gegensatz zu früheren konnektionistischen Modellen enthalten die orthographische und die phonologische Ebene Einheiten, die Grapheme beziehungsweise Phoneme repräsentieren.

Lautes Lesen von Wörtern erfordert nach Plaut et al. (1996) ein orthographisches Muster, welches das passende phonologische Muster generiert. Orthographische und phonologische Informationen interagieren über die Ebene der „hidden units“ solange in beide Richtungen, bis sie den besten Abgleich der jeweiligen Repräsentationsmuster bezüglich eines bestimmten Inputs (Wort bzw. Buchstabenfolge) erreicht haben und die korrekte Aussprache (Output) erfolgen kann. Die Aussprache von Pseudowörtern basiert auf einer Kombination aller bekannter Wortaussprachen, wobei jene, die dem Pseudowort am ähnlichsten sind, den größten Einfluss haben. Entstehen Konflikte bezüglich mehrere Möglichkeiten der Aussprache, werden diese über kooperierende sowie konkurrierende Interaktionen gelöst. (Plaut et al., 1996)

Neben diesem phonologischen Zugangsweg wird auch ein semantischer Zugangsweg angenommen, der die Wortbedeutung miteinbezieht und zusätzlich Informationen aus der Kontextebene einholen kann. Dieser Weg wurde nicht in die computerisierte Form des Triangel-Modells implementiert. Jedoch vertreten die Autoren die Annahme, dass volle Lesekompetenz nur erreicht wird, wenn semantische Einflüsse miteinbezogen werden. (Plaut et al., 1996)

Aus der Sicht verzweigter Netzwerktheorien (Plaut & Gonnerman, 2000) spiegelt Morphologie einen erfolgten Lernprozess wider, der auf der systematischen Zuordnung von oberflächlichen Wortformen (Orthographie, Phonologie) zu deren Bedeutung (Semantik) basiert. Taucht ein bestimmtes orthographisches oder phonologisches Muster in zahlreichen Wörtern auf und stimmt es konsistent mit einer gewissen Bedeutung überein, entsteht eine systeminterne Repräsentation, welche diese Struktur widerspiegelt. Dies führt dazu, dass ein solches Muster im Verarbeitungsprozess als eine einzelne Komponente behandelt wird und somit relativ unabhängig von anderen Teilen des Wortes repräsentiert und verarbeitet wird. Dies gewährleistet, dass neue Kombinationen von bereits bekannten Mustern effizient verarbeitet werden können. (Plaut & Gonnerman, 2000)

Die morphologische Struktur einer Sprache hat einen großen Einfluss auf die Organisation von systeminternen Repräsentationen und den Verarbeitungsprozess. Eine Sprache, die durch eine starke morphologische Struktur³ (z.B.: Hebräisch) geprägt ist, bedingt im Laufe des Lernprozesses, dass die internen Repräsentationen des gesamten Netzwerks stärker morphologisch strukturiert sind. Eine Verarbeitung anhand von Wortstämmen und Affixen wird begünstigt und führt dazu, dass auch Wörter, die eine geringe semantische Transparenz aufweisen, in ihre einzelnen Komponenten zerlegt werden. Ist hingegen die morphologische Struktur einer Sprache schwach ausgeprägt, sind auch die morphologischen Muster im jeweiligen Netzwerk schwächer repräsentiert. Diese weisen daher einen geringeren Einfluss auf die Verarbeitung komplexer, insbesondere semantisch intransparenter Wörter auf. (Plaut & Gonnerman, 2000)

³ Die Struktur der Sprache erlaubt häufig und in hohem Ausmaß, dass Wörter aufgrund ihrer orthographischen/phonologischen Struktur und ihrer Bedeutung systematisch in ihre morphologischen Komponenten gegliedert werden können, die wiederum an der Konstruktion zahlreicher anderer Wörter beteiligt sind (Plaut & Gonnerman, 2000).

2.3 Morpheme als Einheiten der deutschen Schriftsprache

Je nach Sprache beeinflusst Morphologie mehr oder weniger die Produktivität und den Ausdruck von geschriebenen und gesprochenen Wörtern und führt somit zu großen sprachspezifischen Unterschieden (Plaut & Gonnerman, 2000). Das sprachliche Umfeld, in dem ein Kind aufwächst, beeinflusst, welche schriftsprachlichen Kompetenzen ein Kind entwickeln muss, um lesen und schreiben zu lernen (Landerl & Reitsma, 2005).

2.3.1 Definition von Morphemen

Definitionen von Morphemen in der psycholinguistischen Literatur ist gemeinsam, dass die Bedeutung als wichtigste Unterscheidung zu anderen sprachlichen Einheiten (z.B. Silben) erwähnt wird. So werden Morpheme als die kleinsten bedeutungsunterscheidenden (Zwitserslood & Bölte, 2008) oder -tragenden (Arnbak & Elbro, 2000) Einheiten der Sprache bezeichnet. In der schulischen Praxis ist der Begriff „Wortbaustein“ geläufig (Günther, 2010, S. 46, S. 51). Morpheme gelten als fundamentale Bausteine der Wörter in geschriebener und gesprochener Sprache (Kirby et al., 2012). Kuo und Anderson (2006) betonen, dass Morpheme die kleinsten phonologischen Einheiten darstellen, die semantische Informationen tragen. Die Übereinstimmung von semantischen mit phonologischen, nicht mit orthographischen Repräsentationen stellt die Kerneigenschaft von Morphemen dar (Jannedy, Poletto & Weldon, 1994; zitiert nach Kuo & Anderson, 2006).

Grundsätzlich werden in der Orthographie freie („Kind“) von gebundenen Morphemen („-ung“) unterschieden. Zahlreiche Wörter der deutschen Sprache sind aus mehreren Morphemen zusammengesetzt (Zwitserslood & Bölte, 2008). Die Gliederung von zusammengesetzten Wörtern in Morpheme erfolgt anhand des Stammmorphems, welches die lexikalische Bedeutung mit sich trägt, und anhand von gebundenen, grammatischen Morphemen (Affixe), welche die jeweils typischen Rechtschreibregelungen abbilden. Affixe unterscheiden sich je nachdem, ob sie vor (Präfix, „ver-“) oder nach dem Wortstamm (Suffix, „-ung“) angehängt werden. Flexionen bilden Morpheme ab, die Informationen zu Konjugation, Deklination, Steigerungsformen und Pluralform des Wortes bereit stellen („schreib“ = Stamm, „-t“ = Konjugationsmorphem). Suffixe kennzeichnen zudem, ob es sich um ein Adjektiv oder ein Substantiv (Derivation,

„lesbar“, „Lesung“) beziehungsweise eine weibliche Form des Wortes handelt. (Zwitserslood & Bölte, 2008; Schröder-Lenzen, 2013, S. 28f)

Morpheme als kleinste bedeutungstragende Einheiten ermöglichen anhand einer vergleichsweise geringen Bandbreite von phonologischen und orthographischen Einheiten eine enorme Vielfalt von Begriffen auszudrücken und stellen schöpferische Möglichkeiten zur Verfügung, um neue Wörter zu erschaffen (Rastle & Davis, 2008). Die Produktivität eines Morphems gibt an, wie oft dieses Morphem genutzt werden kann, um neue Wörter zu generieren (Elbro & Arnbak, 1996). In der deutschen Sprache finden sich zahlreiche Komposita, die laufend um neue Wortschöpfungen ergänzt werden. Als Komposita werden Wörter bezeichnet, die sich aus selbstständigen Nomen oder/und Adjektiven zusammensetzen („himmelblau“). Manchmal wird ein Fugenmorphem („Engel-s-haar“) hinzugefügt, um die Aussprache zu erleichtern (Zwitserslood & Bölte, 2008).

Morphologisch komplexe Wörter unterscheiden sich in dem Ausmaß ihrer semantischen Transparenz. Wörter werden als transparent bezeichnet, wenn ihre Bedeutung direkt aus der Bedeutung der jeweiligen Morpheme erschlossen werden kann (Plaut & Gonnerman, 2000, Elbro & Arnbak, 1996). Jedoch existieren auch Morpheme, die keine Bedeutung angeben („Him-“ in „Himbeere“) beziehungsweise trotz Gleichschreibung eine andere Bedeutung haben („Schloss“) (Günther, 2010, S.47).

Morpheme sind auf der Basis der Phonographie nicht erfahrbar. Sie können aus Wörtern nicht „herausgehört“ werden, auch beim Sprechen nicht erfahren werden und werden deswegen von Lesanfängern nur schwer identifiziert (Günther, 2010, S. 46, S. 51). Die Silbe stimmt hingegen mit phonologischen Einheiten überein (Günther, 2010, S.46), da sich der natürliche Sprechfluss an der silbischen Struktur orientiert (Schröder-Lenzen, 2013, S.29). Die Beachtung von Morphemen führt zu anderen Segmentierungen als die Orientierung an Silben (Schröder-Lenzen, 2014, S. 29), zum Beispiel: „les-en“ (Trennung in Morpheme) und „le-sen“ (Trennung in Silben). Nach Bredel, Noack und Plag (2011) werden jedoch prosodische (Silben, Akzentsetzung) und morphologische Merkmale (Stamm, gebundene Morpheme) eines Wortes nicht als einander überlagernde oder ausschließende Markierungen in der Wortstruktur gesehen. Vielmehr liefert die Betonung der Silbenstruktur mit dem Trochäus als Basis (auf eine betonte folgt eine unbetonte Silbe) laut Bredel et al. (2011) regelhafte Informationen, wo die Grenze von Stamm und gebundenem

Morphem liegt (zwischen Onset und Nukleus der Reduktionssilbe: „lesen“, „-sen“ ist die Reduktionssilbe, nach dem Onset „s“ liegt die Morphemgrenze).

2.3.2 Prinzip der Morphemkonstanz

In den meisten Sprachen ist das Lernen der Zuordnung von Graphemen zu Phonemen nicht ausreichend, um erfolgreich Lesen und Schreiben zu lernen. Die Aneignung des alphabetischen Prinzips stellt einen wichtigen ersten Schritt im Schriftspracherwerb dar, welcher jedoch durch das Erlernen von sprachspezifischen orthographischen Merkmalen und Besonderheiten ergänzt werden muss (Landerl & Reitsma, 2005). Die deutsche Schriftsprache basiert auf drei grundlegenden Prinzipien, die den Großteil der Wortschreibung bestimmen (Günther, 2010, S. 43):

Das *alphabetische Prinzip*, die phonographische Schreibung, stellt nach Günther (2010, S. 42) die Basis der deutschen Orthographie dar. Zahlreiche Wörter können aufgrund einer eindeutigen Beziehung zwischen Laut und Buchstabe so geschrieben werden, wie sie gesprochen werden.

Das *silbische Prinzip* der Schreibung meint, dass Wörter artikulatorisch und auditiv in Silben gegliedert werden können (Sprechsilbe) und Zugang zu Gleichmäßigkeiten der Schreibweise ermöglichen (Schröder-Lenzen, 2013, S. 27, 32ff). In der deutschen Sprache bauen beispielsweise die meisten Zweisilber auf dem Muster betonte-unbetonte Silbe auf. Folgen dem Vokal der betonten Silbe mindestens zwei Konsonanten, wird der Vokal kurz gesprochen, folgt nur ein oder kein Konsonant, wird der Vokal lang gesprochen, beispielsweise „Betten“ im Gegensatz zu „beten“ (Günther, 2010, S. 42). Folglich kann man aufgrund der Aussprache Rückschlüsse auf die Rechtschreibung ziehen.

Das *morphologische Prinzip* beziehungsweise das Prinzip der Morphemkonstanz führt zu Abweichungen von lautorientierten Schreibweisen und meint, dass sich herkunftsverwandte Wörter, die eine Bedeutung mit sich tragen, in ihrer Schreibweise entsprechen, auch wenn sie unterschiedlich ausgesprochen werden. Beispielsweise entsprechen sich „Hand“ und „Hände“ in ihrer Schreibweise, obwohl „d“ in Hand „t“ ausgesprochen wird und das „ä“ in Hände als „e“. (Günther, 2010, S. 51f)

Das Wissen um morphologische Beziehungen ist für die deutsche Sprache von großer Bedeutung (Landerl & Reitsma, 2005). Besonders bezüglich des Erwerbs von Rechtschreibkenntnissen wird die Bedeutung des Morphems hervorgehoben, da die Kenntnis der Morphemstruktur der deutschen Sprache eine große Rechtschreibhilfe sein kann (Schröder-Lenzen, 2013, S. 28). Wenn die Schreibweise eines Wortstammes bekannt und im Gedächtnis gespeichert ist, kann dieser auch in anderen Wortformen gleich und somit richtig geschrieben werden, vorausgesetzt die Gemeinsamkeit wird erkannt (Günther, 2010, S. 51f; Landerl & Reitsma, 2005). Würde hingegen die morphologische Struktur eines Wortes nicht beachtet, müsste jedes einzelne Wort in seiner jeweiligen Schreibweise im orthographischen Lexikon gespeichert werden (Landerl & Reitsma, 2005). Auch unbekannte komplexe Wörter können anhand der Fähigkeit, sie in ihre Morpheme zu zerlegen, richtig geschrieben werden (Günther, 2010, S. 51f). Da die morphematische Gleichschreibung jedoch nicht durchgängig realisiert wird (Schröder-Lenzen, 2013, S. 28), neigen Kinder bei Ausnahmen, etwa bei starken Verben wie „singen-sangen-gesungen“, zu Überregularisierungen („gesingt“) und können diese Ausnahmen nur durch explizites Erlernen richtig anwenden (Zwitserslood & Bölte, 2008).

Auch im Leseprozess spielt die Morphemorientierung der deutschen Sprache eine bedeutende Rolle. Wenn die Schreibweise eines Morphems im Gedächtnis gespeichert ist, können wiederkehrende Wortbausteine optisch schnell identifiziert werden (Schröder-Lenzen, 2013, 29) und den Leseprozess beschleunigen. Der geübte Leser sollte sowohl silbenorientierte als auch morphemorientierte Strategien nutzen und diese flexibel einsetzen können, je nachdem welche Strukturierungshilfe am effizientesten ist (Schröder-Lenzen, 2013, S.29).

Eine empirische Studie von Landerl und Reitsma (2005) ging der Frage nach, ob Grundschulkinder im deutschen Sprachraum das Prinzip der Morphemkonstanz nutzen, indem die Rechtschreibkenntnisse von Kindern, die holländisch oder deutsch sprechen, verglichen wurden. Im Holländischen wird die Schreibweise von langen Vokalen phonologisch konsistent abgebildet („paar – paren“) im Gegensatz zu der deutschen Schreibweise, die auf einer morphologischen Konsistenz basiert („Paar - Paare“). Die Ergebnisse der Studie zeigen, dass deutschsprachige Kinder bereits in der zweiten Klasse zu 80 Prozent eine konsistente Rechtschreibung für Singular- und Plural-Wortformen verwenden. Dies kann als Indiz gesehen werden, dass die SchülerInnen annehmen,

unterschiedliche Wörter, die den gleichen Wortstamm haben, werden auch gleich geschrieben. Das Erkennen von Gemeinsamkeiten der morphologischen Schreibweise entwickelt sich zunehmend über die Grundschulzeit hinweg und ab der dritten beziehungsweise vierten Klasse scheinen Kinder bereits ein grundlegendes Verständnis für das Prinzip der Morphemkonstanz zu haben. (Landerl & Reitsma, 2005)

2.4 Theoretische Modelle zur Verarbeitung von Morphemen in der Worterkennung

Die Diskussion, wie morphologische komplexe Wörter während dem Lesen beziehungsweise bei der Worterkennung verarbeitet werden, wurde bis dato weder durch theoretische Annahmen noch durch empirische Befunde restlos geklärt. Es stellt sich die Frage, ob komplexe Wörter als eine Einheit aus dem mentalen Lexikon oder als zergliederte Einheiten von Morphemen abgerufen werden. Falls die Verarbeitung von komplexen Wörtern eine Gliederung in Morpheme beansprucht, führt dies erneut zu unterschiedlichen Annahmen, zu welchem Zeitpunkt diese im Prozess der Worterkennung anzuordnen sind. (Bertram, Hyöna & Laine 2011; Beyersmann, Coltheart & Castles, 2012; Bronk, Zwitserlood & Bölte, 2013; Burani, Salmaso & Caramazza, 1984)

Des weiteren stellt sich die grundsätzliche Frage, ob die Verarbeitungsprozesse komplexer Wörter innerhalb eines Systems oder in einem Zwei- beziehungsweise Mehr-Wege-System stattfinden (Bronk et al, 2013). Zusätzlich besteht Uneinigkeit in der Forschung, unter welchen Bedingungen morphologische Einheiten von komplexen Wörtern genutzt werden, ob semantischen oder orthographischen Funktionen von Morphemen eine Bedeutung im Prozess der Gliederung zukommt (Rastle & Davis, 2008; Taft & Nguyen-Hoan, 2010; Bertram et al., 2011).

Grundsätzlich können die Modelle der mentalen Repräsentation anhand folgender Annahmen unterschieden werden (Bronk et al., 2013):

- (1) Morphologisch komplexe Wörter werden bereits vor dem Zugang zum mentalen Lexikon in ihre einzelnen Morpheme zerlegt (Modell der obligatorischen Dekomposition, Taft & Nguyen-Hoan, 2010).

(2) Alle bekannten Wörter werden als Gesamtes im mentalen Lexikon repräsentiert, unabhängig davon ob es sich um komplexe oder einfache Wörter handelt (Butterworth, 1983; zitiert nach Bronk et al., 2013; & nach Zwitserlood & Bölte, 2008).

(3) Morphologische Repräsentationen sind erst zugänglich, nachdem auf die Repräsentation des gesamten Wortes zugegriffen wurde (Supralexikalische Hypothese, Giraudo & Grainger, 2000).

(4) Morphologisch komplexe Wörter können sowohl über Repräsentationen des gesamten Wortes als auch der einzelnen Morpheme verarbeitet werden (Augmented-Addressed-Morphology Model, Burani, Salmaso & Caramazza, 1984; Meta-Modell, Baayen & Schreuder, 1999; Schreuder & Baayen, 1995).

Ad (1): Das Modell der obligatorischen Dekomposition (Taft & Nguyen-Hoan, 2010) geht davon aus, dass Wörter bei der Worterkennung automatisch und zu einem frühen Zeitpunkt des Prozesses in Morpheme zergliedert werden. Ein gegebener Input führt zu einer Aktivierung von orthographischen Einheiten auf der Form-Ebene. Diese orthographischen Einheiten repräsentieren die jeweiligen Morpheme eines Wortstimulus. Die Dekomposition eines komplexen morphologischen Wortes in seine einzelnen Morpheme findet somit bereits vor dem Zugriff auf lexikalische Repräsentationen statt und stellt eine automatische Komponente des Verarbeitungsprozesses visuell präsentierter Wörter dar. Die Aktivierung von orthographischen Repräsentationen wird zu einem sogenannten „Lemma“⁴ übergeleitet. Wörter, die aus einem Morphem zusammengesetzt sind, werden durch ein Lemma repräsentiert, welches als „lexikalischer Eintrag“ des Wortes gesehen werden kann. Aus mehreren Morphemen zusammengesetzte Wörter führen nach ihrer (obligatorischen) Zergliederung zu einer Aktivierung mehrerer Lemmata, welche einen lexikalischen Eintrag des Wortstamms, des Affix sowie des gesamten Wortes repräsentieren. Die Repräsentation des gesamten Wortes gewährleistet, dass die für dieses Wort spezifische Bedeutung zugänglich ist, welche nicht gänzlich durch die Lemmata der einzelnen Morphem-Repräsentationen bestimmt werden kann. Es

⁴ Lemma stellen abstrakte Einheiten dar, die zwischen der (orthographischen) Form- und der übergeordneten (semantischen und syntaktischen) Funktionsebene vermitteln. Im Sinne der netzwerktheoretischen Ansätze können Lemma als stabile Aktivitätsmuster verstanden werden, die sich als Ergebnis einer Korrelation zwischen Input und Output entwickeln. (Taft & Nguyen-Hoan, 2010)

erfolgt keine direkte Aktivierung der Gesamtwort-Repräsentation über die Grapheme der Form-Ebene, sondern erst auf der dazwischenliegenden Ebene der Lemmata. Zusätzlich wird eine Hierarchie von Lemma angenommen, wobei die Repräsentation des Gesamtwortes erst über jene der einzelnen Morpheme aktiviert wird. (Taft & Nguyen-Hoan, 2010)

Rastle und Davis (2008) vertreten basierend auf einer Rezension zahlreicher empirischer Befunde aus Priming-Studien (eine genaue Darstellung siehe Rastle & Davis, 2008) ebenfalls die Annahme, dass im Rahmen des Worterkennungsprozesses jeder geschriebene Stimulus, der eine morphologische Komplexität aufweist, zu einem frühen Zeitpunkt in seine einzelnen Morpheme gegliedert wird. Dieser schnell ablaufende Segmentationsprozess basiert auf einer rein orthographischen Analyse und ist unabhängig von semantischen Informationen des Stimulus. Diese morphologisch-orthographische Dekomposition erlaubt meist einen schnellen Zugang zu der Bedeutung von morphologisch strukturierten Stimuli. (Rastle & Davis, 2008)

Ad (3): Die Supralexikalische Hypothese (Giraudo & Grainger, 2000) beschreibt ein hierarchisches Aktivierungsmodell der Worterkennung, welches Repräsentationen der gesamten Wortform und der zusammensetzenden Morpheme beinhaltet. Im Gegensatz zum oben beschriebenen Modell der obligatorischen Dekomposition sind morphologische Repräsentationen direkt nach Gesamt-Wort-Repräsentationen lokalisiert. Sie vermitteln zwischen dem niedrigeren Level der Formanalyse des Wortstimulus und dem höheren Level der Bedeutung. Alle Wörter, die derselben Wortfamilie angehören, werden über eine gemeinsame morphologische Repräsentation, die dem gemeinsamen Stammmorphem entspricht, verbunden. Die Aktivierung einer morphologischen Einheit ist von der kumulativen Häufigkeit aller Wortrepräsentationen, die dieses Morphem teilen, abhängig. Die Aktivierung einer Wortrepräsentation ist wiederum von der Häufigkeit des Auftretens dieses Wortes in dieser bestimmten Form (Oberflächenfrequenz) abhängig. Somit ist die Aktivierung von morphologischen Repräsentationen eine Funktion der Oberflächenfrequenz des Wortstimulus und führt nur dann zu einer schnelleren Verarbeitung, wenn das gesamte Wort häufig auftritt. (Giraudo & Grainger, 2000)

Ad (4): Das Augmented-Addressed-Morphology-Model (Burani, Salmaso & Caramazza, 1984) geht davon aus, dass während dem Worterkennungsprozess sowohl lexikalische Repräsentationen von Wörtern als auch von Stammmorphemen im mentalen Lexikon

zugänglich sind. Die Annahmen beruhen auf empirischen Studien im italienischen Sprachraum, in denen die Auswirkung der Häufigkeit des Auftretens von Wörtern beziehungsweise von Stammmorphemen auf lexikalische Entscheidungsaufgaben untersucht wurde. Hochfrequente, bekannte Wörter können über den direkten Zugang zu der Einheit des gesamten Wortes effizienter und schneller verarbeitet werden. Unbekannte oder selten auftretende Wörter werden in Abhängigkeit der kumulativen Häufigkeiten aller morphologisch verwandten Wörter (Frequenz des Stammmorphems) verarbeitet. Das heißt, dass Wörter, die als Gesamtes selten auftreten oder unbekannt sind, deren Wortstamm jedoch häufig auftritt, schneller abgerufen werden können als Wörter, die ebenso selten als Gesamtes auftreten oder unbekannt sind und zugleich auf einem niederfrequenten Wortstamm beruhen. Somit werden zwei Wege der Verarbeitung angenommen, der direkte lexikalische Zugang zur Wortrepräsentation als auch der zu morphologischen Repräsentationen nach Gliederung des Wortes in seine Bestandteile, sofern es sich um ein selten auftretendes oder neues Wort handelt (Burani et al., 1984; Caramazza, Laudanna & Romani, 1988, zitiert nach Bronk et al., 2013).

Ad (4): Das Meta-Modell der morphologischen Verarbeitung (Schreuder & Baayen, 1995; Baayen & Schreuder, 1999) ähnelt dem Augmented-Addressed-Morphology-Model (AAM-Modell, Burani et al., 1984) in der Annahme zweier möglicher Verarbeitungsrouten von morphologisch komplexen Wörtern. Ein wesentlicher Unterschied liegt jedoch in der Annahme, dass im Gegensatz zum AAM-Modell keine strikte Trennung bezüglich der Verarbeitung von bekannten und unbekannten komplexen Wörtern getroffen wird. Vielmehr können bei jedem komplexen Wort beide Routen aktiviert werden, welche miteinander in Interaktion stehen und um den lexikalischen Zugriff auf die Repräsentation des gesamten Wortes oder der jeweiligen Morpheme konkurrieren („lexical competition“). Im Verarbeitungsprozess verläuft die Aktivierung von Morphemen über drei Stufen. Die erste Stufe betrifft die Zuordnung eines sprachlichen Inputs zu formbasierten Zugriffsrepräsentationen⁵, welche gesamte Wörter als auch Wortstämme und Affixe abbilden können. Je besser die Zuordnung eines Inputs zu einer Zugriffsrepräsentation gelingt, desto stärker steigt deren Aktivitätsniveau an und führt gleichzeitig zu einem Senken der Aktivierung anderer Zugriffsrepräsentationen. Das Erreichen eines Schwellenwerts führt zur Aktivierung

⁵ Die Häufigkeit des Auftretens des jeweiligen sprachlichen Inputs bestimmt das individuelle Ruheaktivitätsniveau der zugehörigen Zugriffsrepräsentation (Schreuder & Baayen, 1995).

zugehöriger Repräsentationen und leitet die zweite Stufe ein. In diesem Stadium werden die Übereinstimmung von gemeinsam aktivierten morphologischen Repräsentationen und die Angemessenheit ihrer Kombinationen durch Miteinbeziehen von semantischen und syntaktischen Informationen überprüft. Nach erfolgreicher Lizenzierung werden schließlich die lexikalischen Repräsentationen der Komponenten integriert und zu einer Repräsentation des gesamten Wortes zusammengefügt. Parallel zu dieser beschriebenen Route läuft die direkte Route über den Zugriff auf die Repräsentation des gesamten Wortes. Die Geschwindigkeit und Effizienz der jeweiligen Route wirkt sich rückwirkend auf die zu Beginn aktivierten lexikalischen Einheiten aus. Führt zum Beispiel die Aktivierung eines Morphems zu einem schnelleren Ergebnis, wird dessen Repräsentation gestärkt, indem das Ruheaktivitätsniveau erhöht wird. Somit kann zukünftige Aktivierung dieser morphologischen Repräsentation schneller einen Schwellenwert erreichen und weitergeleitet werden. Dasselbe gilt für gesamte Wortrepräsentationen. Die semantische Transparenz⁶ und die Häufigkeit des Auftretens des gesamten Wortes oder seiner zusammensetzenden Morpheme sind ausschlaggebend für den Beitrag der jeweiligen Routen für den Prozess der Worterkennung von komplexen Wörtern. (Baayen & Schreuder, 1999; Schreuder & Baayen, 1995)

Der Frage, welches dieser geschilderten Modelle am ehesten die Verarbeitung morphologisch komplexer Wörter abbilden kann, wurde in einer experimentellen Studie von Beyersmann, Coltheart und Castles (2012) nachgegangen. Unterschiedliche Versuchsbedingungen wurden geschaffen, um zu erheben, ob morphologisch komplexe Wörter im Verarbeitungsprozess immer zuerst einer Gliederung in einzelne Morpheme unterlaufen oder ob diese über den Zugriff auf eine Repräsentation des gesamten Wortes verarbeitet werden. In rein lexikalischen Entscheidungsaufgaben und in solchen mit vorherigem maskierten Priming wurde der Effekt von vertauschten Buchstaben⁷ auf morphologisch komplexe Wörter untersucht. In einer weiteren Primingbedingung wurde

⁶ Komplexe Wörter, die keine semantische Transparenz aufweisen (deren einzelnen Morpheme nicht die Bedeutung des gesamten Wortes widerspiegeln), werden schneller über den Zugriff auf Repräsentationen des gesamten Wortes verarbeitet. Hingegen können semantisch transparente Wörter (Bedeutung des gesamten Wortes wird durch seine einzelnen Morpheme widergespiegelt), die in dieser Kombination selten auftreten, schneller über die Route des Parsen verarbeitet werden. (Schreuder & Baayen, 1995)

⁷ Ein Wort dessen Buchstaben vertauscht wurden („drak“) führt aufgrund einer ungenauen Dekodierung in frühen Stadien des Worterkennungsprozesses zu einer Aktivierung des ähnlichsten lexikalischen Eintrages („dark“). Im Worterkennungsprozess versucht das System nun den ursprünglichen Wortstimulus mit diesem Eintrag übereinzustimmen. In lexikalischen Entscheidungsaufgaben (Entscheidung, ob es sich um ein existierendes Wort oder ein Pseudowort handelt) kann dies zu einer Verlangsamung der Antwort führen. (Beyersmann et al., 2012)

die semantische Transparenz des Primes variiert. Zusammenfassend bieten die empirischen Ergebnisse dieser Studie Hinweise für beide Formen der Verarbeitung, sowohl über den direkten Zugriff auf eine lexikalische Repräsentation des gesamten Wortes als auch über den Weg der morphologischen Dekomposition. Laut den Autoren der Studie lassen sich diese Ergebnisse am ehesten im Sinne des Meta-Modells (Schreuder & Baayen, 1995) erklären, welches eine parallele, simultane Aktivierung von Repräsentationen des gesamten Wortes und der jeweiligen Morpheme annimmt, um dem Leser eine möglichst effiziente Worterkennung zu gewährleisten. Diese Ergebnisse beschränken sich jedoch auf den englischen Sprachraum. (Beyersmann et al., 2012)

Eine Studie von Bronk, Zwitserlood und Bölte (2013) im deutschen Sprachraum zeigt, dass in lexikalischen Entscheidungsaufgaben komplexe Wörter schneller bearbeitet werden, wenn sie sich aus zwei freien Morphemen zusammensetzen („Papierhut“), als wenn sie nur aus einem Morphem bestehen („Margerite“). Die Voraussetzung dafür ist, dass jedes einzelne Morphem für sich betrachtet eine höhere Frequenz des Auftretens aufweist („Papier“, „Hut“) als das gesamte Wort. Dann ermöglicht die Aufspaltung des Wortes einen schnelleren und effizienteren Zugang zum mentalen Lexikon als der direkte Zugriff auf eine Einheit des gesamten Wortes. Diese Beobachtung konnte auch bei semantisch-intransparenten Komposita („Spaßvogel“) gemacht werden und führt zu der Annahme, dass eine Dekomposition von Wörtern in die jeweiligen Morpheme bereits vor dem Zugriff auf semantische Informationen stattfindet. Die volle Bedeutung eines zergliederten Wortes kann jedoch nur über ein erneutes Zusammenfügen der einzelnen Teile erlangt werden. Semantisch transparente Komposita erleichtern aufgrund der ähnlichen Bedeutung ihrer Morpheme zum gesamten Wort („Baumhaus“) die Worterkennung. Semantisch-intransparente Wörter können zu längeren Verarbeitungsprozessen führen, da ihre Teile nicht der Bedeutung des gesamten Wortes entsprechen. Die Ergebnisse dieser Studie lassen sich laut den Autoren am ehesten anhand des Modells von Taft (2010) erklären, welches eine frühe Dekomposition von Wörtern in morphologische Einheiten und den Zugriff auf die Bedeutung des Wortes mittels übergeordneter, zusammenfügender Lemmas erklärt. (Bronk et al, 2013)

Die meisten Studien zur Verarbeitung morphologisch komplexer Wörter umfassen Priming-Aufgaben, da mit Hilfe dieser Aufgaben Aussagen getroffen werden können, ob

eine morphologische Ebene in die visuelle Worterkennung involviert ist und wann diese im Verarbeitungsprozess auftritt. Mittels der Vorgabe eines morphologischen Primes im Gegensatz zu einem rein orthographischen oder zusammenhangslosen Prime lässt sich feststellen, ob dieser zu einem schnelleren Erkennen eines Zielwortes führt. Durch Manipulation der Dauer der Primevorgabe lassen sich Aussagen darüber treffen, zu welchem Zeitpunkt morphologische Effekte im Verarbeitungsprozess auftauchen. (Casalis, Dusautoir, Colé & Ducrot, 2009)

In einer empirischen Priming-Studie an französischen SchülerInnen der vierten Schulstufe untersuchten Casalis et al. (2009), ob Kinder in der Worterkennung für die morphologische Struktur eines geschriebenen Wortes empfänglich sind. Morphologisches Priming führte zu einer schnelleren lexikalischen Entscheidung bezüglich des Zielwortes („mariage“, Heirat) als zusammenhangsloses Priming. Eine kurze Vorgabe des Primes führte zu vergleichbaren Effekten von morphologischen („marier“, heiraten) als auch von orthographischen Primes („marine“). In einer längeren Primingbedingung führten nur morphologische Primes zu einer kürzeren Latenzzeit. Somit unterscheidet sich die morphologische Aktivierung zu einem früheren Zeitpunkt im Worterkennungsprozess nicht signifikant von einer orthographischen, jedoch im weiteren Verlauf (bei längerer Dauer des Primevorgabe). Die Ergebnisse dieser Studie zeigen, dass morphologische Verarbeitung in den Prozess der Worterkennung involviert ist, und weisen auf unterschiedliche Ebenen von Repräsentationen hin. (Casalis et al., 2009)

2.5 Entwicklung des Wissens um morphologische Einheiten

In den geschilderten Modellen wird angenommen, dass dem Leser im Worterkennungsprozess (zu welchem Zeitpunkt beziehungsweise unter welchen Bedingungen auch immer) morphologisch strukturierte Repräsentationen zur Verfügung stehen. Da jedoch ein geschriebenes Wort nicht von sich aus visuelle Hinweise zur Verfügung stellt, anhand welcher orthographischer Anhaltspunkte Morphem-Grenzen identifiziert werden können, stellt sich die Frage, wie Leser im Laufe der Entwicklung lernen, welche Buchstabeneinheiten im geschriebenen Text morphologische Einheiten abbilden (Rastle & Davis, 2008).

Laut Rastle und Davis (2008) kann der Leser die statistische Struktur von Buchstabensequenzen im Rahmen der Aneignung von morphologisch-orthographischem Wissen nutzen. Eine Analyse der Auftretenswahrscheinlichkeiten von Buchstabenkombinationen im geschriebenen Text erlaubt einerseits die Trennung von selten auftretenden Buchstabensequenzen, welche dadurch gleichzeitig eine Morphemgrenze markiert (z.B.: das in der englischen Sprache wenig wahrscheinliche Buchstabenpaar „pf“ markiert in dem Wort „helpful“ die Morphem-Grenze „help-ful“). Andererseits besteht die Möglichkeit häufig auftretende Buchstabenkombinationen zu einzelnen morphologischen Einheiten zusammenzufassen. Entscheidend für effizientes Zusammenfassen von Buchstabenmustern ist nicht nur deren hohe Auftretenswahrscheinlichkeit, sondern auch die Beachtung mit welchen anderen linguistischen Einheiten sie in Kombination auftreten (z.B. Affixe, die mit bestimmten Wortstämmen gemeinsam auftreten). (Rastle & Davis, 2008)

Eine weitere mögliche Hypothese nach Rastle und Davis (2008) besagt, dass das Erlernen von morphologischen Einheiten aufgrund einer Zuordnung der orthographischen Form zur Bedeutung passiert. Beginnenden Lesern steht bereits ein großer Wortschatz an gesprochenen Wörtern zu Verfügung, der lexikalische und semantische Repräsentationen von üblichen Wortstämmen und Affixen der jeweiligen Sprache umfasst. Im Laufe des Schriftspracherwerbs lernt der Leser, dass bestimmte schriftliche Buchstabensequenzen morphologische Einheiten darstellen, die bereits aus der gesprochenen Sprache bekannt sind. In zahlreichen unterschiedlichen Wörtern, die Wortstamm oder Affixe teilen, können Leser konsistente Zusammenhänge von Schreibweise und Bedeutung entdecken. Dies führt nach Giraudo und Grainger (2000) im Laufe des Schriftspracherwerbs zu einer zunehmenden Festigung von morphologischen Repräsentationen im Gedächtnis, welche durch die explizite Instruktion, dass morphologische Beziehungen zwischen Wörtern bestehen, erleichtert werden kann. Auch die Häufigkeit des Auftretens eines Morphems beeinflusst dessen Speicherung als mentale Repräsentation. Ein Wortstamm, der sowohl selbstständig als auch mit gebundenen Morphemen kombiniert häufig auftritt und einer großen Wortfamilie angehört, wird in seiner mentalen Repräsentation gestärkt und kann schließlich auch in einer neuen, unbekannten Zusammensetzung leichter und schneller erkannt werden (Carlisle & Stone, 2005).

Im Verlauf der Leseentwicklung wird der Entwicklung des Wissens um morphologische Strukturen eine unterschiedliche Bedeutung zugesprochen. Nach Arnbak und Elbro (2000) lassen sich grundsätzlich drei theoretische Annahmen über die Beziehung von Lesefertigkeiten und morphologischer Bewusstheit unterscheiden. Bereits junge Kinder verfügen über eine Art intuitives Wissen über die Morphologie der (gesprochenen) Sprache (Berko, 1959; Clark & Hecht, 1982; Tornéus, 1987; zitiert nach Arnbak & Elbro, 2000), welches sich auf die spätere Entwicklung von Lesefertigkeiten auswirkt. Im Gegensatz dazu steht die Annahme, dass sich morphologische Bewusstheit durch die Erfahrung mit der geschriebenen Sprache entwickelt, da dadurch erst ermöglicht wird, ein und dasselbe Morphem in vielen unterschiedlichen Wörtern erkennen zu können. Die dritte Annahme besagt, dass morphologische Bewusstheit und Leseentwicklung über einen dritten zugrundeliegenden Faktor verbunden sind, die phonologische Bewusstheit, und vielleicht nur ein scheinbarer Zusammenhang zwischen den beiden ersteren besteht. (Arnbak & Elbro, 2000)

Entwicklungsstudien weisen darauf hin, dass die morphologische Bewusstheit stark mit der phonologischen Bewusstheit zusammenhängt, und führen zu der Vermutung, dass beide eine gemeinsame Fähigkeit abbilden könnten. Die Größe dieses Zusammenhangs ist in Hinblick auf die morphologische Bewusstheit von Menschen mit Lese-Rechtschreib-Schwäche (LRS) bedeutsam. Die meisten Kinder mit LRS weisen große phonologische Defizite auf, welche bei bestehendem Zusammenhang auch die Entwicklung von morphologischen Fähigkeiten beeinträchtigen müssten. Im Gegensatz dazu steht die Option, dass sich morphologische Bewusstheit unabhängig von der phonologischen entwickelt und auf dem Vermitteln von semantischen Einheiten in der gesprochenen Sprache beruht. (Casalis, Colé & Soppo, 2004)

Eine Studie von Casalis et al. (2004) an französischen Kindern ging der Frage nach, ob eine Beeinträchtigung der phonologischen Bewusstheit bei LRS die Verarbeitung von größeren Einheiten, wie Morphemen, beeinflusst. Bei Aufgabestellungen, eine phonologische Abfolge in Morpheme zu segmentieren, zeigten Kinder mit LRS stärkere Beeinträchtigungen als jüngere, der Altersnorm entsprechende Leser. Dies deutet einen Zusammenhang, zwischen der Fertigkeit Phoneme und Morpheme zu erkennen, an. In Aufgabestellungen, die produktives Wissen darüber erforderten, wie Wörter abgeleitet werden, zeigten Kinder mit LRS jedoch keine speziellen Probleme. Demnach entwickeln

sich morphologische Fertigkeiten zumindest teilweise unabhängig von phonologischen Fertigkeiten. Kinder mit LRS können sich eher auf semantische als auf phonologische Informationen verlassen und morphologische Fertigkeiten als kompensatorische Strategien nutzen, um schwach ausgeprägte phonologische Fertigkeiten zu umgehen. (Casalis et al., 2004)

Laut Kuo und Anderson (2006) werden in der Leseforschung zahlreiche Interpretationen des Begriffs „morphologische Bewusstheit“ verwendet und es sollte ihrer Ansicht nach zwischen der Aneignung der Morphologie und der Entwicklung der morphologischen Bewusstheit unterschieden werden. Ersteres bezeichnet die Entwicklung der Fähigkeit morphologisch komplexe Wörter in der natürlichen Sprache zu verstehen und zu produzieren. Letzteres konzentriert sich auf das Erlangen der Fähigkeit, Regeln der Wortbildung, die mögliche Kombinationen von Morphemen bestimmen, auch unabhängig von einem kommunikativen Kontext zu reflektieren und zu manipulieren. Das Wissen um morphologische Regeln kann als ein Teilgebiet der erworbenen Kenntnisse über Morphologie gesehen werden. (Kuo & Anderson, 2006)

Studien in unterschiedlichen Schriftsprachen zeigen, dass morphologische Bewusstheit mit fortschreitender Leseentwicklung zunehmend an Bedeutung gewinnt. Leseanfänger sind mit der Bewältigung von graphophonischen Analysen beschäftigt. Sobald das Wissen um die Zuordnung von Graphem zu Phonem und umgekehrt gefestigt ist und die Anforderungen in Texten durch morphologisch komplexe Wörter steigen, erlangt die morphologische Bewusstheit mit Verlauf der Grundschulzeit immer mehr an Wichtigkeit (Kuo & Anderson, 2006). Mit Hilfe bekannter Morpheme kann die Bedeutung eines neuen komplexen Wortes erfasst werden (Kuo & Anderson, 2006) und die zunehmende Leseerfahrung begünstigt wiederum den Aufbau von mentalen morphologischen Repräsentationen (Carlisle & Stone, 2005; Kuo & Anderson, 2006).

Ein umfassendes Verständnis, wie sich die morphologische Bewusstheit und die Verarbeitung von Morphemen bei Kindern entwickelt, ist für den Schriftspracherwerb von Bedeutung und sollte sowohl in Forschungen zur Entwicklung des Lesens Beachtung finden (Casalis et al., 2004; Rastle & Davis, 2008), als auch in Modelle der Leseentwicklung miteinbezogen werden (Carlisle & Stone, 2005). Das Erkennen und Verarbeiten von Morphemen kann einen wichtigen Beitrag zu schnellem und effizientem Lesen leisten und sollte in den Methoden der (schulischen) Leseinstruktion stärker präsent

sein (Rastle & Davis, 2008). In der Literatur finden sich pädagogische Empfehlungen, dass morphologisch basierte Instruktion erst erfolgen sollte, wenn Kinder das Wissen über orthographische Muster bereits gründlich erworben und automatisiert haben (Adam, 1990; zitiert nach Kuo & Anderson, 2006). Dem entgegen steht die Annahme, dass Kinder bereits in frühen Grundschulstufen von einer expliziten Instruktion profitieren können, die das Erkennen von morphologischen Strukturen in geschriebener und gesprochener Sprache fördert (Carlisle & Stone, 2005; Kuo & Anderson, 2006). Morpheme als bedeutungstragende Einheiten können das Einprägen von orthographischen Mustern begünstigen (Carlisle, 2003; zitiert nach Kuo & Anderson, 2006), das Lesen unbekannter Wörter erleichtern und zur Entwicklung des Wortschatzes beitragen (Carlisle & Stone, 2005). Zukünftige Forschung ist notwendig, um der Frage nachzugehen, zu welchem Zeitpunkt das Wissen über (morphologische) Wortstrukturen in den Unterricht eingebaut werden soll (Kuo & Anderson, 2006).

2.6 Verarbeitung von Morphemen beim lauten Lesen

Bezüglich der Verarbeitung von Morphemen existieren fast keine Studien, in denen ihre Bedeutung beim lauten Lesen von Kindern untersucht wird. Die meisten Studien umfassen lexikalische Entscheidungsaufgaben, meist mit vorangegangenem Priming. Der Prozess des lauten Lesens unterscheidet sich jedoch von Aufgaben zur lexikalischen Entscheidung, welche notwendigerweise eine semantische Ebene miteinbeziehen, aber nicht das Generieren der Aussprache des präsentierten Wortes erfordern. Folglich können die unterschiedlichen Aufgabestellungen verschiedene Hinweise auf bedeutsame Komponenten des Leseprozess geben.

Studien im englischen Sprachraum zeigen, dass morphologische Strukturen beim Lesen von Wörtern bereits bei GrundschülerInnen eine Rolle spielen. In einer Studie von Carlisles und Stone (2006) wurde untersucht, ob bekannte Morpheme das Lesen von Wörtern erleichtern und beschleunigen. SchülerInnen niedriger und hoher Grundschulstufen wurden phonologisch und orthographisch transparent abgeleitete⁸, zweisilbige Wörter präsentiert, die sich entweder aus einem Stamm und einem Suffix

⁸ Transparent abgeleitet meint, dass sich Schreibweise und Aussprache des Wortstamms („hill“) nicht durch eine Derivation zu einem Adjektiv („hilly“) verändern (Carlisle & Stone, 2006).

(„hilly“) oder nur einem Morphem („silly“) zusammensetzten. Die Wörter wurden sowohl in Hinblick auf ihre Wortfrequenz als auch auf ihre Wortlänge und Schreibweise gepaart. Der Unterschied zwischen den Wörtern lag somit allein in deren morphologischer Struktur, welche nur eine Teilung der erstgenannten Wortstimuli im Leseprozess erlaubte. Sowohl die SchülerInnen unterer als auch höherer Grundschulstufen lasen Wörter, die sich aus zwei Morphemen zusammensetzten, schneller und genauer. Die Ergebnisse weisen darauf hin, dass das Erkennen eines häufigen Wortstammes den Leseprozess erleichtert. (Carlisle & Stone, 2006)

Bereits im vorherigen Kapitel wurde angeführt, dass das Erkennen von morphologischen Einheiten im Leseprozess das Benennen von komplexen Wörtern, insbesondere jenen, die in dieser Kombination neu sind, erleichtern kann (Carlisle & Stone, 2006; Rastle & Davis, 2005). Burani, Marcolini und Stella (2002) spezifizieren diese Annahme auf den Prozess des lauten Lesens, in welchem sich die Verfügbarkeit von Wortstämmen und Suffixen positiv auf die Lesegeschwindigkeit und -genauigkeit von neuen komplexen Wörtern auswirkt, da eine Zuordnung von Graphem zu Phonem einen arbeitsintensiveren und somit längeren Prozess bedeuten würde. In einer transparenten Orthographie, welche eine verlässliche Zuordnung von Graphem zu Phonem erlaubt, wird meist angenommen, dass Kinder zu Beginn der Leseentwicklung auf diese Strategie zurückgreifen, um neue Wörter zu benennen, und erst im weiteren Verlauf der Leseentwicklung die morphologische Struktur der Sprache nutzen. Jedoch existieren sehr wenige Studien in transparenten Orthographien, die morphologische Effekte beim lauten Lesen von Kindern erfassen. (Burani et al., 2002)

Die Studie von Burani, Marcolini und Stella (2002) ist eine der ersten Studien an Grundschulkindern in einer transparenten Orthographie (Italienisch), welche die morphologische Verarbeitung nicht allein anhand lexikalischer Entscheidungsaufgaben, sondern auch bei Aufgaben zum lauten Lesen untersuchte. Das Stimulusmaterial umfasste Pseudowörter, die sich aus Morphemen zusammensetzten, jedoch in dieser Kombination nicht existierten („donnista“), und Pseudowörter, die davon abgeleitet wurden, jedoch keine tatsächlichen Morpheme enthielten („dennosto“). Wenn das mentale Lexikon von Kindern in Morpheme strukturiert ist, sollten Pseudowörter aus echten Morphemen schneller und genauer gelesen werden als solche, die keine Morpheme enthalten. Eine Graphem-Phonem Zuordnung, wie sie bei neuen, unbekannten Wörtern ohne

morphologische Struktur erwartet wird, würde einen längeren Verarbeitungsprozess bedeuten. Eine zusätzliche Annahme besagt, dass, wenn die morphembasierte Verarbeitung mit Alter und Lesekompetenz wächst, Kinder in höheren Klassen vermehrt Wörter anhand ihrer morphologischen Strukturen erkennen sollten. Die Ergebnisse der Studie zeigen, dass die Kinder aller Altersgruppen und Schulstufen (8-10 Jahre; 3-5. Schulstufe) Pseudowörter aus Morphemen schneller und genauer lasen als die davon abgeleiteten Pseudowörter. Des Weiteren konnte in Aufgaben zu lexikalischen Entscheidungen festgestellt werden, dass Kinder aller Schulstufen häufiger dazu tendierten, Pseudowörter aus tatsächlichen Morphemen (fälschlicherweise) als mögliche Wörter zu beurteilen. Diese Ergebnisse können als Hinweis auf eine frühe Organisation des mentalen Lexikons in Stämme und Suffixe gesehen werden und lassen vermuten, dass bei Aufgaben zum lauten Lesen eine parallele Aktivierung der semantischen Komponente, welche die Bedeutung der Morphem-Kombination überprüft, nicht unbedingt notwendig ist. (Burani et al., 2002)

Eine weitere Studie von Burani, Marcolini, De Luca & Zoccolotti (2008) bezog zusätzlich zu unterschiedlicher Altersstufen von Kindern auch deren unterschiedlichen Lesefähigkeiten, einschließlich Kinder mit Leserechtschreibschwäche (LRS), mit ein. Ähnlich zu der Studie von Burani et al. (2002) konnte aufgezeigt werden, dass alle Kinder unabhängig von ihrem Alter und ihren Lesefertigkeiten Pseudowörter, die sich aus Morphemen zusammensetzten, schneller und genauer lasen. Zusätzlich wurden Wörter präsentiert, die in Morpheme segmentiert werden konnten („cass-iere“), und Wörter von ähnlicher Länge und Häufigkeit, die jedoch nicht getrennt werden konnten („cammello“). Geübte Leser zeigten in dieser Versuchsbedingung keine signifikanten Unterschiede im lauten Lesen bezüglich Geschwindigkeit und Genauigkeit. Im Gegensatz dazu konnten Kinder mit LRS Wörter, die sich in Morpheme teilen ließen, signifikant schneller und genauer laut vorlesen als Wörter gleicher Länge, die jedoch keine Trennung in Stamm und Suffix erlaubten. Das Morphem kann Kindern mit geringeren Lesefertigkeiten (LRS, jüngere Kinder) als eine unterstützende Einheit von mittlerer Größe dienen: ganze Wörter sind eine zu große Einheit und reine Graphem-Phonem-Zuordnungen beanspruchen einen genauen, langsamen Analyseprozess. Morphembasiertes Lesen stellt demnach eine effiziente Lesestrategie für schwache Leser dar. (Burani et al., 2008)

In einer weiteren Studie von Marcolini, Traficante, Zoccolotti & Burani (2011) wurde

zusätzlich zu unterschiedlichen Fähigkeiten der Leser auch die Wortfrequenz miteinbezogen. Die Ergebnisse zeigen, dass Wörter, die in Morpheme zerlegt werden konnten, von schwachen Lesern schneller gelesen wurden, als Wörter, die nicht getrennt werden konnten, und zwar unabhängig von deren Häufigkeit. Dieser Effekt zeigte sich bei jungen, geübten Lesern jedoch nur dann, wenn es sich um seltene Wörter handelte. Häufige, aus Morphemen zusammengesetzte wurden vergleichbar schnell gelesen wie häufige Wörter, die nicht in morphologische Einheiten zerlegt werden konnten. Somit profitieren nur schwache Leser auch bei sehr häufigen Wörtern von der Möglichkeit diese in morphologische Einheiten zu gliedern. (Marcolini et al., 2011)

Nach Burani et al. (2008) und Marcolini et al. (2011) findet morphologische Verarbeitung beim lauten Lesen statt, wenn der Gewinn höher als die Kosten ist. Das Trennen eines bekannten Wortes in seine Morpheme kann zu Kosten des Verarbeitungsprozesses gehen und zu einer Verlangsamung im Prozess des lauten Lesens führen, wenn die Aussprache des Wortes generiert wird. So kann sich die Akzentsetzung in einem Wortstamm verändern, wenn ihm eine zuvor abgetrennte Endung erneut hinzugefügt wird, und eine erneute Generierung der Aussprache dieser Kombination erfordern. Wenn hingegen ein lexikalischer Gesamtwortseintrag zum schriftlichen Wort nicht verfügbar oder noch nicht ausreichend gefestigt ist, wie bei neuen oder unbekannten Wörtern, führt eine Verarbeitung von Morphemen zu einer rascheren Aussprache, als wenn einzelne Grapheme zu Phonemen zugeordnet werden müssen. (Marcolini et al., 2011)

Die beschriebenen Studien im italienischen Sprachraum zeigen, dass Leseanfänger und schwache Leser über den möglichen Zugriff auf morphologische Einheiten ihre Lesegeschwindigkeit erhöhen (Marcolini et al. 2011). Eine Studie von Elbro und Arnbak (1996) in einer weniger transparenten Orthographie (Dänisch) brachte vergleichbare Resultate. Jugendliche mit diagnostizierter Leserechtschreibschwäche profitierten von der morphologischen Struktur vorgegebener Wörter. Morphologisch transparente Wörter („sunburn“) wurden schneller und genauer laut gelesen als solche, die keine transparente Struktur („window“) aufwiesen. Jüngere, geübte Leser zeigten keine Unterschiede in der Bearbeitung der jeweiligen Wortarten. Des Weiteren konnte gezeigt werden, dass leseschwache Jugendliche, die jedoch die morphologische Analyse in der Worterkennung besonders nutzten, zu den besten Lesern hinsichtlich ihres Leseverständnisses gehörten. Nach Elbro und Arnbak (1996) liegt für schwache Leser der Vorteil einer Dekomposition eines Wortes in seine Morpheme in dem Zugriff auf die semantische Ebene des Wortes

über seine (bedeutungstragenden) Morpheme. Die Verarbeitung von Morphemen stellt bei Leserechtschreibschwäche demnach über den Zugriff auf semantische Informationen eine kompensatorische Strategie in der Worterkennung und im Wortverständnis dar und ermöglicht Defizite im phonologischen Rekodieren auszugleichen (Ebro & Arnbak, 1996). Diese Annahme findet sich auch in der bereits beschriebenen Studie von Casalis et al. (2004) wieder.

Die Autoren der zuvor beschriebenen Studie im italienischen Sprachraum (Burani et al., 2008) gelangen jedoch aufgrund des fehlenden Einflusses der semantischen Transparenz auf die Lesegeschwindigkeit zu einem anderweitigen Schluss. Bei Aufgaben zum lauten Lesen unter Zeitdruck ist eine möglichst schnelle Zuordnung von einer korrekten Artikulation zu dem geschriebenen Wort notwendig, welche zumindest in transparenten Orthographien auch ohne Zugriff auf die semantische Ebene möglich ist. Morpheme werden als eine phonologisch und orthografisch verlässliche Leseinheit für schwache Leser gesehen. (Burani et al. 2008)

Wie bereits in der Einleitung erwähnt, können empirische Befunde aus anderen Sprachen nicht unmittelbar auf den deutschen Sprachraum übertragen werden. Neben sprachlich bedingten Unterschieden kann auch die im jeweiligen Schulsystem erfolgte Instruktion Einfluss auf die „bevorzugte“ Verwendung von bestimmten orthographischen Einheiten nehmen (siehe dazu auch Bedeutung der Instruktion nach Klicpera et al., 2010). Die vorliegende Arbeit geht in einer im Folgenden beschriebenen Studie der Frage nach, welche Verwendung Morpheme im Prozess der lauten Lesens bei österreichischen Schülerinnen, Ende der Grundschule, finden, und soll einen Beitrag zu der Fragestellung leisten, welche Einheiten Kindern im Rahmen der Leseentwicklung zur Verfügung stehen und von ihnen effizient genutzt werden können.

3 Empirie

In diesem Kapitel werden die Entstehung der Daten, die methodische Vorgehensweise, die Erhebung, die durchgeführten Analysen und die daraus folgenden Ergebnisse dargestellt.

3.1 Fragestellung und Hypothesen

Ausgehend von der Forschungsfrage „Welche Bedeutung haben Morpheme für das laute Lesen bei Kindern am Ende der vierten Klasse Grundschule?“ wurden Leseaufgaben konstruiert, die eruieren sollen, ob Morpheme als größere Einheiten während dem lauten Lesen verfügbar sind. Insbesondere soll der Frage nachgegangen werden, wie diese im Vergleich zu anderen Einheiten, wie Silben oder Graphemen, die Worterkennung unterstützen und die Geschwindigkeit des Leseprozesses beeinflussen.

Das eigens für diese Untersuchung konzipierte Design (genaue Darstellung in Kapitel 3.2.2 Design der Untersuchung S.45) beinhaltet zwei unterschiedliche Herangehensweisen, eine Pseudowortbedingung und eine Trennbedingung, um der Frage nachzugehen, welche Verwendung beziehungsweise Bedeutung Morpheme beim lauten Lesen von Kindern, Ende der vierten Klasse Volksschule, finden.

Folgende Hypothesen sollen überprüft werden:

Hypothese der Pseudowortbedingung: „Kinder lesen Pseudowörter, die aus existierenden Morphemen zusammengesetzt sind, schneller laut vor als Pseudowörter, die keine morphologischen Einheiten enthalten.“

Hypothese der Trennbedingung: „Kinder lesen Wörter, deren Morphemstrukturen intakt sind, schneller laut vor als Wörter, deren Morphemstrukturen gestört sind.“

Die zugrundeliegende Annahme der Hypothesen ist, dass die Lesegeschwindigkeit im Sinne der Zwei-Wege-Modelle Aussagen über das Maß der Verfügbarkeit eines lexikalischen Zugangs zu morphologischen Repräsentationen zulässt. Auch im Sinne der Netzwerkmodelle wird die Lesegeschwindigkeit als Maß der Aktivierung von

morphologischen Mustern gewertet und soll demnach Rückschlüsse auf die Effizienz von Morphemen im Leseprozess erlauben.

Wenn größere Einheiten, im vorliegenden Fall morphologische Einheiten, verfügbar sind, müsste demnach gemäß der den Hypothesen zugrundeliegenden Annahme ein Pseudowort, welches sich aus Morphemen zusammensetzt, schneller laut gelesen werden als ein vollkommen „sinnfreies“ Pseudowort, welches theoretisch nur über (zeitintensive, da aufwendigere) Graphem-Phonem-Zuordnung rekodiert werden kann.

Die Trennung in Silben oder in Morpheme geht der Frage nach, welche dieser Einheiten im gegenseitigen Vergleich zu einer höheren Lesegeschwindigkeit führt. Sollten Morpheme schneller aus dem Gedächtnis abrufbar sein beziehungsweise zu einer stärkeren Aktivierung von Verbindungen führen, würden die SchülerInnen Wörter, deren Morpheme intakt bleiben (bei gleichzeitiger Störung der Silbenstruktur) schneller laut vorlesen als Wörter, deren Silben aufrechterhalten bleiben (bei gleichzeitiger Störung der Morphemstruktur).

3.2 Methode

3.2.1 Beschreibung der TeilnehmerInnen

Nach dem Einholen des Einverständnisses zur Durchführung der Studie von Landesschulrat, DirektorInnen und LehrerInnen wurde die Datenerhebung Ende des Schuljahres an zwei Volksschulen im oberösterreichischen Innviertel durchgeführt. Insgesamt konnten 43 SchülerInnen der vierten Schulstufe mit einer vorab schriftlich bestätigten Einverständniserklärung der Erziehungsberechtigten teilnehmen. Fünf der erfolgten Sprachaufnahmen der Kinder mussten jedoch aufgrund technischer Mängel oder massiv störender Hintergrundgeräusche komplett aus der Analyse ausgeschlossen werden. Somit setzte sich die in die Auswertung miteinbezogene Stichprobe letztlich aus 38 SchülerInnen, 16 Mädchen und 22 Buben, die zum Zeitpunkt der Testung im Durchschnitt zehn Jahre und vier Monate alt waren, zusammen.

3.2.2 Design der Untersuchung

Als Erhebungsinstrument der vorliegenden Studie dienten zwei eigens für diese Untersuchung konzipierte Aufgabestellungen, welche im Folgenden näher beschrieben werden.

3.2.2.1 Leseaufgabe: Pseudowörter

Angelehnt an die Studie von Burani, Marcolini und Stella (2002) wurden Paare von Pseudowörtern, Wörter, die gelesen und ausgesprochen werden können, jedoch an sich keine Bedeutung tragen, konstruiert. Jedes Paar umfasste ein Pseudowort, welches aus realen Morphemen, je einem Stamm und einem Suffix, besteht, jedoch in der gegebenen Zusammensetzung im tatsächlichen Sprachgebrauch nicht vorkommt. Parallel dazu wurde ein Pseudowort entwickelt, welches weder ein existierendes Stammmorphem noch ein tatsächliches Suffix enthält, jedoch in Wortaufbau und -bild dem Ursprungsmorphem vergleichbar bleibt. Erzielt wurde dies durch Ersetzen der Vokale im ursprünglichen Morphem mit neuen Vokalen unter Beibehalten der vorhandenen Konsonanten.

Das folgende Beispiel soll dem besseren Verständnis der Vorgehensweise dienen: Aus dem Wort „sagen“ wurde das Stammmorphem „sag-“ entnommen und mit dem Suffix „-ig“ verbunden, somit entsteht das Pseudowort „sagig“. Durch ein Ersetzen der Vokale mit anderen entsteht „sugog“, wobei weder „sug-“ noch „-og“ existierende Morpheme darstellen.

Auf diese Weise wurden zehn Paare von Pseudowörtern, somit insgesamt 20 Stimuli konstruiert (siehe Anhang A). Bei der Auswahl der Wörter, von denen die Pseudowörter abgeleitet wurden, wurde mittels CELEX-Datenbank⁹ (Baayen, Piepenbrock & van Rijn, 1993) darauf geachtet, dass diese über eine vergleichbare Auftretenshäufigkeit verfügen (2–4 log/million). Des Weiteren wurden die ursprünglichen Wörter mit Unterrichtsmaterialien zum Grundwortschatz von VolksschülerInnen abgeglichen.

⁹ CELEX stellt eine computerbasierte lexikalische Datenbank zur Verfügung, welche Informationen zu Orthographie, Phonologie, Morphologie, Syntax und Frequenz für sechs Millionen deutsche Wortinträge, die aus dem Mannheimer Korpus des Instituts für deutsche Sprache generiert wurden, bereitstellt (German Linguistic Guide by Gulikers, Rattink, Piepenbrock, 1993).

3.2.2.2 Leseaufgabe: Störung

Für die zweite Aufgabe wurden Wörter in drei unterschiedlichen Trennbedingungen präsentiert. Eine Trennung beziehungsweise Störung des Wortes wurde mittels des Zeichens „#“ erzeugt, da dieses zwar keinem Buchstaben ähnelt, somit als Störung identifiziert werden kann, sich jedoch ins Schriftbild eingliedern lässt.

1. Die erste Bedingung umfasste eine Trennung mittels „#“ in Morpheme („sing#en“). Diese kommt einem Aufrechterhalten des Stammmorphems und des Suffix gleich bei gleichzeitiger Störung der Silben-Struktur.
2. In der zweiten Bedingung wurde das Zeichen an die Silbengrenzen gesetzt, somit das Wort in Silben getrennt bei Störung der Morphem-Struktur („sin#gen“).
3. Die dritte Bedingung stellte keine Trennung im eigentlichen Sinne dar. Das gesamte Wort blieb intakt, jedoch wurde vor beziehungsweise nach dem Wort ein Raute-Zeichen eingefügt („#singen“).

Um die unterschiedlichen Trennbedingungen trotz der geplanten Durchführung der Erhebung zu einem einzigen Zeitpunkt miteinander vergleichen zu können, wurden zehn Wortgruppen von je drei Wörtern gebildet (entsprechend der Anzahl der Trennbedingungen). Innerhalb einer Gruppe wurden die Wortlänge und die Abfolge von Konsonanten und Vokalen konstant gehalten. Auch das Auftreten von Doppelkonsonanten wurde berücksichtigt. Mittels CELEX-Datenbank (Baayen et al., 1993) wurde sichergestellt, dass die Wörter innerhalb einer Gruppe eine ähnliche Häufigkeit aufwiesen (siehe Anhang B). Zusätzlich wurde ihr Auftreten im Wortschatz von Volksschülern anhand gängiger Unterrichtsmaterialien abgeglichen.

Jedes der insgesamt 30 Wörter durchlief jede Trennbedingung, somit entstanden drei Blöcke mit je 30 unterschiedlichen Wörtern. In jedem Block blieben zehn Wörter intakt und je zehn Wörter wurden an den Silben- beziehungsweise Morphem-Grenzen getrennt (siehe Anhang B).

3.2.3 Durchführung

Aus den gewonnenen Stimuli der Leseaufgabe „Störung“ ergaben sich entsprechend der drei Blöcke drei Präsentationen. Innerhalb einer Darbietung kam jedes der drei Wörter der zehn Wortgruppen in einer anderen Trennbedingung vor. Da drei Präsentationen konstruiert wurden, konnte auch jedes Wort der Dreier-Wortgruppen in jeder Trennbedingung dargestellt werden. Die 30 Wörter pro Block wurden zufällig gemischt dargeboten.

Die Pseudowort-Paare wurden auf zwei Blöcke aufgeteilt, wobei jeder Block das jeweilige Gegenstück des Paares enthielt. Jeder Block bestand aus fünf Pseudowörtern, die sich aus existierenden Morphemen zusammensetzten, und fünf gänzlich sinnfreien Pseudowörtern. Die Blöcke wurden entsprechend der drei Präsentationen zufällig gemischt und jeweils ein Block vor, der andere Block mit den entsprechenden Gegenständen nach der Leseaufgabe „Störung“ dargeboten.

Folglich ergaben sich Listen von 50 Wörtern, welche in Schriftgröße 75 der Schriftart Century Gothic mittels Powerpoint-Präsentation auf einem Laptop mit 15,4 Zoll Bildschirmdiagonale vorgegeben wurden. Die Aufnahme erfolgte mittels Diktiergerät. Mit Hilfe eines Audioeditor-Programms wurden die gewonnen Tonspuren zu einem späteren Zeitpunkt ausgewertet und kodiert.

Bevor ein Stimuli am Bildschirm erschien, wurde ein großes „X“ als Orientierungspunkt gezeigt (siehe folgende Abbildung).

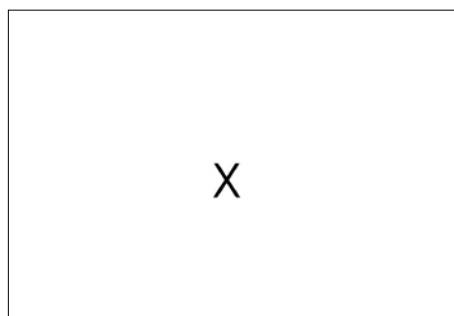


Abbildung 1: Fixierungs- bzw. Orientierungspunkt am Bildschirm

Beim darauf folgenden Erscheinen eines Wortes (siehe folgende Abbildung) erfolgte zeitgleich ein kurzer, markanter Ton (ein Klick), um zum einen die Aufmerksamkeit der TeilnehmerInnen zu sichern und zum anderen eine spätere Auswertung der Tonspur zu ermöglichen. Der Klick zeigte den Beginn der benötigten Lesezeit an.



Abbildung 2: Präsentation des Stimulus am Bildschirm

Jedes Wort wurde vom Kind laut gelesen. Das jeweils präsentierte Wort wurde solange gezeigt, bis das Kind den Leseprozess erkennbar abgeschlossen hatte. Erst dann ging der Testleiter per Mausklick zum nächsten Wort weiter.

Die Datenerhebung an den Schulen fand während der regulären Unterrichtszeit statt und wurde als Einzeltestung in einem von der Schule zur Verfügung gestellten, ruhigen und separaten Raum durchgeführt. Jedes Kind erhielt zu Beginn der Präsentation die gleichbleibende Instruktion, die folgenden Wörter so schnell und so genau wie möglich laut vorzulesen. Des Weiteren wurden die SchülerInnen zu Beginn der Leseaufgabe „Pseudowörter“ darauf hingewiesen, dass nun Wörter auftauchen würden, die keine „echten“ Wörter seien, jedoch so wie diese gelesen werden könnten. Bevor die eigentliche Testung startete, wurden den Kindern drei Übungswörter vorgegeben. Nach dem ersten Teil der Testung wurden die Kinder wieder anhand von drei Übungswörtern mit den Trennbedingungen vertraut und darauf aufmerksam gemacht, dass sich in die Wörter ein Zeichen „geschummelt“ hätte, welches sie jedoch nicht beachten sollten. Bei dem darauf folgenden Block der parallelisierten Pseudowörter wurden sie an die Instruktion zu Beginn erinnert.

3.2.4 Analyse und Kodierung der Tonspuren

Nach der Durchsicht mehrerer Tonspuren wurde die Gesamtlesezeit der Latenzzeit¹⁰ für die weitere Analyse vorgezogen, da die SchülerInnen unterschiedliche Strategien beim Vorlesen zeigten. Einige Kinder begannen beispielsweise fast unmittelbar nach der Präsentation des Wortes mit der Aussprache und lasen das Wort in einem Zug laut vor. Diese kurze Latenzzeit führte somit zu einer vergleichsweise kurzen Gesamtlesezeit des präsentierten Wortes. Jedoch waren bei zahlreichen Schülern Herangehensweisen zur Bewältigung der Aufgabenstellung, die Wörter laut vorzulesen, beobachtbar, welche eine gleichwertige Interpretation der Latenzzeiten nicht erlaubt hätten. Manche Kinder begannen sofort nach der Präsentation ein Wort auszusprechen (kurze Latenzzeit), wobei sie dieses erst während des bereits begonnen Leseprozesses zu erlesen schienen (lange Gesamtlesezeit). Im Gegensatz dazu benötigten einige Kinder vergleichsweise länger bis sie mit der Aussprache begannen (lange Latenzzeit), lasen dann jedoch das präsentierte Wort in einem Zug vor (verhältnismäßig kürzere Gesamtlesezeit). Auch mehrmaliges Ansetzen bis ein Wort gelesen wurde, war beobachtbar. Die Gesamtlesezeit entsprach somit in diesen geschilderten Fällen einem verlässlicheren Maß, um Unterschiede in der Lesegeschwindigkeit in Abhängigkeit von den präsentierten Wortarten je Bedingung zu erfassen.

Die präsentierten Wörter und Pseudowörter wurden entsprechend der Lesegenauigkeit als „spontan richtig“, „spontan falsch“ (ohne Korrektur) oder „selbst korrigiert“ (zuerst falsch, jedoch richtig nach Selbstkorrektur) kodiert. Auch Wortstimuli, bei denen Kinder die Betonung nicht vollkommen passend setzten, wurden als richtig gewertet. Jedoch wurden die „tatsächlichen“ Wörter der Trennbedingung, welche durch eine übermäßig falsche Betonung zu einem eindeutig falschen Ergebnis führten (z.B.: „Hasen“ vs. „hassen“) auch als solches kodiert. Als „ungültig“ wurde kodiert, wenn die Lesezeit eines Wortes aufgrund von Hintergrundgeräuschen, Aussetzen des Klicks oder andauerndem Lachen der Kinder während der Aussprache nicht eindeutig erhoben werden konnte.

¹⁰ Latenzzeit meint die Zeit ab Beginn der Wortpräsentation bis zum Beginn der Aussprache durch das Kind. Die Gesamtlesezeit entspricht der Zeit ab Beginn der Wortpräsentation bis zum Ende der Aussprache des Wortes.

3.3 Ergebnisse

Der Grunddatensatz der vorliegenden Untersuchung, in welchem jedes einzelne Item, also jedes vorgegebene (Pseudo-)Wort aufgeschlüsselt wurde, besteht aus 1900 Datensätzen. Dieser Grunddatensatz wurde zwei getrennten einfaktoriellen Varianzanalysen unterzogen, um zum einen den Einfluss des Faktors Geschlechts und zum anderen den Einfluss der besuchten Schulklasse auf die Lesezeit der präsentierten Wörter allgemein und somit ohne Berücksichtigung der jeweiligen Wortbedingung zu erheben. Die Voraussetzung der Homogenität der Varianzen wurde in beiden Analysen mittels Levene-Test überprüft (Faktor Geschlecht: $p=.140$; Faktor Schulklasse: $p=.302$). Die folgenden Ergebnisse werden unter der Annahme eines Signifikanzniveaus von $\alpha = .05$ berichtet.

Bezüglich der unabhängigen Variable Geschlecht zeigte sich kein signifikanter Unterschied zwischen Mädchen und Buben in Bezug auf die Lesezeit, $F(1, 1837) = .390$, $p=.532$. Die präsentierten Items (Pseudowörter und Wörter) werden somit von Mädchen und Buben gleichermaßen schnell vorgelesen, wie die durchschnittlichen Lesezeiten der folgenden Tabelle zeigen:

	Mittelwert	Standardabweichung	N
männlich	1,480	,541	1061
weiblich	1,496	,592	778
Gesamt	1,487	,563	1839

Tabelle 1: Mittlere Lesegesamtzeiten (in Sekunden, auf Tausendstel gerundet) nach Geschlecht, Standardabweichung, Anzahl der Items (unter Ausschluss der ungültigen Lesezeiten)

Bezüglich der drei unterschiedlichen Schulklassen konnte ein systematischer Unterschied (mittels ANOVA) in Hinblick auf die Lesezeit festgestellt werden, $F(2, 1836)=3.196$, $p=.041$. In einer Post Hoc Analyse mittels Tukey HSD (honest significant difference) Test wurden die Lesezeiten pro Schulklassen mehrfach verglichen und es zeigte sich, dass sich nicht alle drei Schulklassen übergreifend bezüglich der benötigten Lesezeit voneinander unterscheiden, sondern allein der Unterschied zwischen zwei Klassen signifikant wird ($p=.041$). In der unten stehenden Tabelle werden die durchschnittlichen Lesezeiten der jeweiligen Schulklasse dargestellt. Die mittleren Lesezeiten der Klasse 2 fallen am

niedrigsten aus. Die präsentierten (Pseudo-)Wörter wurden von diesen SchülerInnen signifikant schneller gelesen als von jenen der Klasse 1. Die Lesezeiten der Klasse 3 weisen zu keiner der beiden Klassen signifikante Unterschiede auf (zu Klasse 1: $p=.952$, zu Klasse 2: $p=.143$).

	Mittelwert	Standardabweichung	N
Klasse 1	1,513	,559	783
Klasse 2	1,437	,555	564
Klasse 3	1,503	,577	492
Gesamt	1,487	,563	1839

Tabelle 2: Mittlere Lesegesamtzeiten (in Sekunden, auf Tausendstel gerundet) nach besuchter Klasse, Standardabweichung, Anzahl der Items (unter Ausschluss der ungültigen Lesezeiten)

Nach genauer Analyse der Tonspuren wurde entschieden, dass nur spontan richtig gelesene Wörter in die weitere Auswertung einbezogen werden, da auch bei selbst korrigierten Wörtern mögliche Versuchsleitereffekte nicht ausgeschlossen werden können (siehe dazu auch Diskussion der Methode S.60). Für jedes Kind wurde für jede der insgesamt fünf Lesebedingungen (zwei Pseudowortbedingungen, drei Trennbedingungen) ein arithmetischer Durchschnittswert der Gesamtlesezeit gebildet. Somit wurde der ursprüngliche Datensatz zu einem Datensatz auf Basis der Kinder verdichtet. Im Folgenden werden die Ergebnisse analog zum Design der Studie anhand der zwei unterschiedlichen Versuchsbedingungen dargestellt.

3.3.1 Ergebnisse der Pseudowortbedingungen

Die durchschnittlichen Gesamtlesezeiten der Kinder pro Pseudowortart werden im ersten Schritt mittels deskriptiver Kennwerte dargestellt. Wie folgende Tabelle zeigt, unterscheiden sich die mittleren Gesamtlesezeiten je nach Art des Pseudowortes. Pseudowörter, die aus tatsächlichen Morphemen zusammengesetzt wurden („könnung“), weisen eine geringere durchschnittliche Lesezeit auf im Vergleich zu Pseudowörtern, die mittels Ersetzen der Vokale (siehe auch 3.2.2.1 Leseaufgabe: Pseudowörter, S.45) davon abgeleitet wurden („kännang“).

	Mittelwert	Standardabweichung	N
Pseudowort aus 2 echten Morphemen	1,602	,416	38
Pseudowort von 2 Morphemen abgeleitet	1,811	,414	38

Tabelle 3: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Pseudowortbedingung, Standardabweichung und Anzahl der SchülerInnen

t-Test: Eine inferenzstatistische Analyse erfolgte mittels t-Test für abhängige Stichproben. Die vorausgesetzte Normalverteilung der Differenzen der beiden Pseudowort-Lesezeiten pro Kind wurde mittels Kolmogorov-Smirnov-Test überprüft und kann als gegeben angenommen werden, $p=.183$ (mit Korrektur nach Lilliefors). Alle folgenden Ergebnisse werden unter der Annahme eines Signifikanzniveaus von $\alpha=.05$ berichtet. Für die Gesamtlesezeiten zeigte sich, dass Pseudowörter, die aus tatsächlichen Morphemen konstruiert wurden, von den SchülerInnen hoch signifikant schneller laut vorgelesen wurden als Pseudowörter, die davon abgeleitet wurden und keine existierenden Morpheme enthielten, $t(37)=-4.470$, $p<.001$. Des Weiteren korrelieren die Gesamtlesezeiten der beiden Pseudowortbedingungen, $r=.785$, $p<.001$. Somit zeigten die einzelnen SchülerInnen eine ähnliche Leseperformanz in den jeweiligen Pseudowortbedingungen. Kinder, die Pseudowörter aus „echten“ Morphemen schneller lesen, lesen auch solche aus „unechten“ Morphemen schneller.

Varianzanalyse mit Messwiederholung: Um den Einfluss der unabhängigen Variablen Geschlecht sowie Zugehörigkeit zu der jeweiligen Schulklasse auf die Lesezeiten der jeweiligen Pseudowortbedingungen sowie etwaige Wechselwirkungen betrachten zu können, wurde zusätzlich eine Varianzanalyse mit Messwiederholung durchgeführt. In dem folgenden statistischen Modell dienten als Innersubjektfaktoren die zwei unterschiedlichen Pseudowortarten und als Zwischensubjektfaktoren die jeweiligen drei Schulklassen und das Geschlecht der SchülerInnen.

Ein signifikanter Haupteffekt der Art der Pseudowortbedingung auf die Lesezeit der SchülerInnen bestätigt den zuvor mittels t-Test beschriebenen Unterschied in den Lesezeiten, $F(1,32)=19.612$, $p<.001$. Des Weiteren konnten keine signifikanten Wechselwirkungen zwischen den Lesezeiten der jeweiligen Pseudowortbedingung und dem Geschlecht beziehungsweise der Schulklasse festgestellt werden. Die durchschnittlichen Gesamtlesezeiten von Mädchen als auch Buben wurden ähnlich durch die Pseudowortbedingungen beeinflusst, $F(1,32)=2,407$, $p=.131$, wie auch die Mittelwerte der folgenden Tabelle zeigen.

	Männlich (N=22)		Weiblich (N=16)	
	M	SD	M	S
Pseudowort aus 2 echten Morphemen	1,611	,447	1,589	,382
Pseudowort von 2 Morphemen abgeleitet	1,755	,384	1,887	,454

Tabelle 4: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Pseudowortbedingung nach Geschlecht geteilt, Standardabweichung und Anzahl der SchülerInnen

Jedoch zeichnete sich ein Einfluss der Variable Schulklasse auf die Lesezeit der Pseudowörter ab, $F(2,32)=2.736$; $p=.08^{11}$. Die in nachstehender Tabelle dargestellten mittleren Lesezeiten der SchülerInnen weisen bei Pseudowörtern, die aus tatsächlichen Morphemen konstruiert wurden, in allen drei Schulklassen ähnliche Werte auf ($M_1=1.5836$, $M_2=1.6359$, $M_3=1.5889$). Jedoch unterscheidet sich bei Pseudowörtern, welche keine existierenden Morphemen enthalten, die mittlere Lesezeit der Klasse 1 ($M_1=1.9217$) stärker von jenen der anderen beiden Klassen ($M_2=1.7540$, $M_3=1.7007$) und weist somit die höchste Differenz zwischen den beiden Bedingungen (Pseudowort aus „echten“ Morphemen und Pseudowort aus „falschen“ Morphemen) auf.

¹¹ Bei einer Varianzanalyse mit Messwiederholung unter Ausschluss des Faktors Geschlecht erhöht sich der Anteil der erklärten Varianz des Faktors Schulklasse von 14,6% auf 16,8% und das Ergebnis wird signifikant $F(2,35) = 3.541$, $p=.040$.

	Klasse 1 (N=16)		Klasse 2 (N=12)		Klasse 3 (N=10)	
	M	S	M	S	M	S
Pseudowort aus 2 echten Morphemen	1,584	,318	1,636	,484	1,589	,504
Pseudowort von 2 Morphemen abgeleitet	1,922	,517	1,754	,329	1,701	,295

Tabelle 5: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Pseudowortbedingung nach Klasse, Standardabweichung und Anzahl der SchülerInnen

Die nachstehende graphische Darstellung veranschaulicht die beschriebenen Differenzen in der Leseperformanz je Klasse:

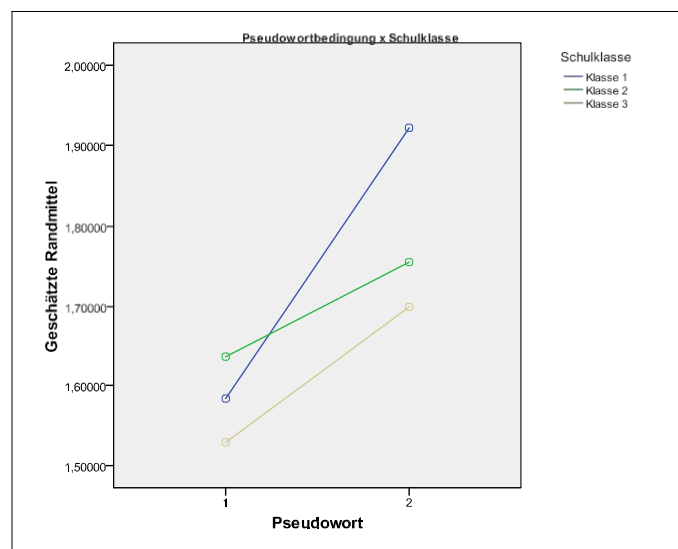


Abbildung 3: Profildiagramm der geschätzten Randmittelwerte für die Gesamtlesezeiten je Pseudowortbedingung (1=Pseudowort aus Morphemen, 2=abgeleitetes Pseudowort ohne Morpheme) und Schulklasse.

Die SchülerInnen aller drei Klasse zeigten eine vergleichbare Leseperformanz bei Pseudowörtern aus Morphemen, jedoch erhöhte sich die Lesezeit bei Kindern der Klasse 1 überproportional, wenn Pseudowörter ohne jegliche morphologische Struktur gelesen wurden (siehe blaue Verbindungslinie mit deutlich höherem Anstieg).

3.3.2 Ergebnisse der Trennbedingungen

Die in folgender Tabelle angeführten deskriptiven Kennwerte zeigen einen ersten Überblick der durchschnittlichen Gesamtlesezeiten je nach Art der Trennung innerhalb eines Wortes. Wörter, die keinerlei Trennung erfuhren („lesen#“), weisen die kürzeste mittlere Gesamtlesezeit auf. Wörter, die an der Morphemgrenze getrennt wurden und deren morphologische Strukturen somit aufrechterhalten blieben („les#en“), wurden durchschnittlich schneller laut vorgelesen als solche, die an der Silbengrenze getrennt wurden und deren Morphem-Strukturen dadurch gestört wurden („le#sen“).

	Mittelwert	Standardabweichung	N
Trennung an Morphemgrenze	1,283	,186	38
Trennung an Silbengrenze	1,327	,172	38
Keine Trennung	1,220	,153	38

Tabelle 6: Mittlere Gesamtlesezeiten (in Sekunden, auf Tausendstel gerundet) der jeweiligen Trennbedingung, Standardabweichung und Anzahl der SchülerInnen

t-Test: Nachdem eine Normalverteilung der Differenzwerte der Wörter mit Trennung an der Morphemgrenze und an der Silbengrenze angenommen werden kann (Überprüfung mittels Kolmogorov-Smirnov-Test mit Korrektur nach Lilliefors, $p=.200$), wurde für die Gesamtlesezeiten dieser beiden Trennbedingungen ein t-Test für abhängige Stichproben berechnet. Die Abhängigkeit ergibt sich aus dem speziellen Versuchsaufbau dieser Leseaufgabe (siehe dazu Kapitel 3.2.2.2 Leseaufgabe: Störung S.46).

Die SchülerInnen wiesen bei Wörtern, die an der Morphemgrenze getrennt wurden, eine signifikant kürzere Lesezeit auf als bei Wörtern, die an der Silbengrenze getrennt wurden, $t(37)=-2.667$, $p=.011$. Somit wurden Wörter, deren morphologische Strukturen intakt blieben, schneller gelesen als Wörter, deren Morpheme gestört wurden.

Varianzanalyse mit Messwiederholung: Die mittleren Gesamtlesezeiten aller drei Trennbedingungen wurden im weiteren Verlauf der Auswertung einer Varianzanalyse mit Messwiederholung unterzogen, welche aufgrund der bereits beim t-Test erwähnten Konstruktion dieser Leseaufgabe (siehe dazu S.46) möglich war. Diese Analyse erlaubt zusätzliche Aussagen über etwaige Wechselwirkungen sowie Einflüsse der unabhängigen Variablen Geschlecht und Schulklasse. In dem im Folgenden beschriebenen Modell dienten als Innersubjektfaktoren die drei Trennbedingungen und als Zwischensubjektfaktoren das Geschlecht der Kinder sowie die besuchte Schulklasse.

Eine Überprüfung der Sphärizität mittels Mauchly-Test zeigte, dass diese Voraussetzung (gerade) nicht verletzt wurde, $\chi^2(2)=5.761$, $p=.056$. Somit konnte auf eine Korrektur der Freiheitsgrade verzichtet werden. Es erwies sich ein hoch signifikanter Haupteffekt der Trennbedingung der Wörter auf die Gesamtlesezeit der SchülerInnen, $F(2,64)=21,460$, $p<.001$. Die Art der Worttrennung wirkte sich somit auf die benötigte Lesezeit der Kinder aus.

Paarweise Vergleiche lassen erkennen, dass sich alle drei Bedingungen signifikant voneinander unterscheiden. Wie auch die durchschnittlichen Gesamtlesezeiten pro Trennbedingungen in Tabelle 6 zeigen, wurden Wörter, die keinerlei Trennung erfuhren, von den SchülerInnen schneller laut vorgelesen als die Wörter der beiden Trennbedingungen¹². Wie im zuvor durchgeführten t-Test unterscheiden sich auch die mittleren Lesezeiten der beiden Trennbedingungen signifikant, $p=0.031$. Wörter, die an der Morphemgrenze getrennt wurden, wiesen niedrigere Gesamtlesezeiten auf als solche, die an der Silbengrenze getrennt wurden. Somit wurden Wörter, deren morphologische Strukturen erhalten blieben, von den SchülerInnen schneller laut vorgelesen als Wörter, deren Morpheme gestört wurden. Die folgende Abbildung veranschaulicht, dass Wörter ohne Trennung verhältnismäßig am schnellsten laut vorgelesen wurden, gefolgt von Wörtern, deren Morphemstruktur erhalten war. Wörter, deren Silbenstruktur erhalten blieb und deren Morpheme somit gestört waren, wurden von den Schülern vergleichsweise am langsamsten gelesen.

¹² Der Unterschied von Lesezeiten der Wörter ohne Trennung wird zu jenen mit Trennung an der Morphemgrenze mit $p=0.001$ und zu jenen mit Trennung an der Silbengrenze mit $p<0.001$ signifikant.

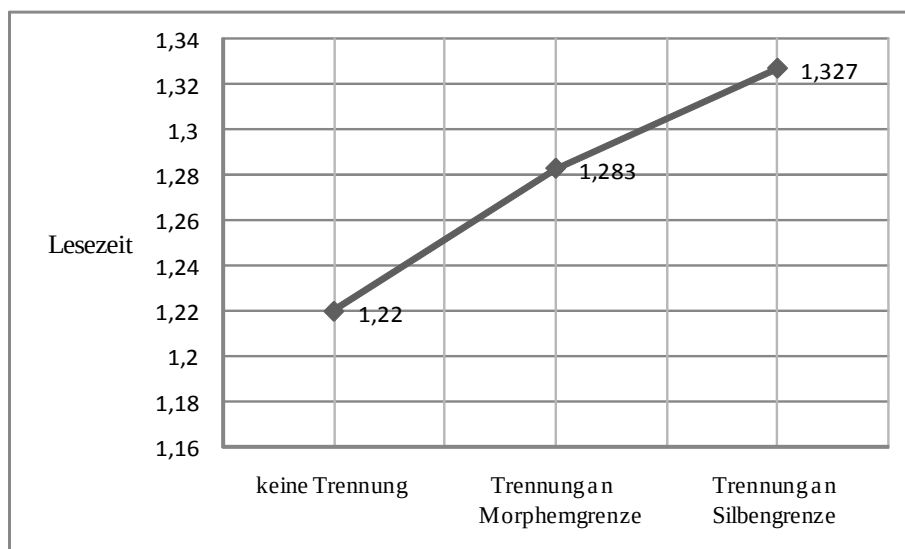


Abbildung 4: graphische Darstellung der mittleren Lesezeiten je Trennbedingung in Sekunden

Die unabhängigen Variablen Geschlecht und Schulklasse zeigten keinerlei Einflüsse auf die abhängige Variable der Gesamtlesezeit je nach Trennbedingung. Somit wurden Buben und Mädchen gleichermaßen durch die jeweiligen Trennbedingungen beeinflusst, $F(2, 31)=.887$, $p=.422$. Des Weiteren zeigten die SchülerInnen über alle drei Klassen hinweg ähnliche Lesezeiten je Trennbedingung, $F(4, 64)=1.540$, $p=.201$.

4 Diskussion

4.1 Methode

Die Wortliste der Erhebung umfasste 50 Wörter und erwies sich angesichts der Testung von SchülerInnen während der Unterrichtszeit als angemessen. Eine längere Wortliste wäre mit einem größeren Zeitaufwand sowie der Frage der Zumutbarkeit verbunden. Bei der Auswahl der Stimuluswörter wurde eine mittlere und, insbesondere innerhalb der Wortgruppen der Trennbedingung, vergleichbare Worthäufigkeit mittels CELEX-Datenbank (Baayen et al., 1993) sichergestellt. Jedoch sollte in zukünftigen Studien auch die Frequenz der jeweiligen Einheiten (Morpheme, Silben) beachtet und konstant gehalten werden. Ein beabsichtigtes Variieren der Morphem- und Silbenfrequenz in zukünftigen Studien könnte zusätzlich Aufschlüsse bieten, ob dies die Verarbeitung der jeweiligen Einheit beeinflusst (z.B. begünstigt oder hemmt) und welche Auswirkung dies auf den Prozess des lauten Lesens zeigt.

In zukünftigen Erhebungen an Kindern wäre es wünschenswert, wenn eine Datenbank zur Verfügung steht, welche die Frequenz von Wörtern, Silben und Morphemen im Wortschatz von Kindern, idealerweise auch je Altersklasse, beachtet. Auch wenn in der vorliegenden Arbeit die verwendeten Wörter mit gängigen Unterrichtsmaterialien abgeglichen wurden, kann nicht automatisch angenommen werden, dass sich diese Worte mit der von der Datenbank CELEX angegebenen logarithmierten Worthäufigkeit ebenso im Sprachgebrauch von Kindern finden lassen.

Die Items (Wörter und Pseudowörter) dienten unmittelbar nach deren Konstruktion als Erhebungsinstrument bei den teilnehmenden SchülerInnen. Ein Testlauf im Vorfeld wurde im Rahmen der vorliegenden Arbeit nicht durchgeführt. Eine Vorstudie ist jedoch empfehlenswert für zukünftige Forschungen, da dadurch mögliche Effekte, welche sich aufgrund der Itemkonstruktion ergeben könnten, geprüft und in der Hauptstudie kontrolliert werden können.

Während der Durchführung der vorliegenden Studie wurden mögliche Effekte sichtbar, welche im Folgenden beschrieben werden und durch deren Beachtung zukünftige Studien optimiert werden können:

(1) In der Pseudowortbedingung wurde unter anderem zwischen den Endungen „-ig“ („sagig“) und „-ung“ („könnung“) variiert. Die Endung „-ung“ kennzeichnet in der deutschen Sprache ein Nomen, jedoch wurde bei der Vorgabe der Wörter alles in kleingeschriebener Schrift präsentiert. Es kann somit nicht ausgeschlossen werden, dass ein Ignorieren der „Pseudowort-Grammatik“ durch Vernachlässigung der üblichen Großschreibung einen störenden beziehungsweise zusätzlich verzögernden Einfluss auf die Worterkennung hatte.

(2) Auch die Aussprache eines Wortes kann einen Effekt auf die Worterkennung haben. So wurde beispielsweise das Wort „holer“ wiederholt als „Holler“ ausgesprochen, ein Wort, das es tatsächlich gibt. In diesen Fällen kann nicht eindeutig auf eine Aktivierung von den Morphemen „hol-“ und „-er“ im Gegensatz zu dem Gesamtworteintrag „Holler“ geschlossen werden, zumindest auf Stufe der Artikulation.

(3) Innerhalb der Wortgruppen der Trennbedingung wurden Häufigkeit und Vokal-Konsonanten-Abfolge konstant gehalten, jedoch wurde zwischen Verben, Pronomen und Nomen variiert. Die konstant gehaltenen Wortheigenschaften wurden als besonders wichtig für den Leseprozess erachtet und somit auch auf Kosten der dadurch entstandenen Variation bezüglich der Wortart berücksichtigt.

(4) Auch die Wortstruktur kann sich auf die Worterkennung auswirken und die Verarbeitung von bestimmten Einheiten begünstigen oder verzögern. Beispielsweise führt eine Trennung in Silben zu einer Trennung von Doppelkonsonanten („ret#ten“) und kann unter Umständen eine zusätzliche Auswirkung auf den Leseprozess bedeuten.

(5) Des Weiteren können in der Trennbedingung mögliche Priming-Effekte nicht ausgeschlossen werden. So wurde beispielsweise das Wort „rennen“ präsentiert und nach einigen Wörtern „kennen“. Somit bleibt die Frage offen, ob das vorherige Lesen des einen Wortes das Erkennen des Folgenden erleichtert.

Neben den beschriebenen möglichen Effekten, welche auf die Konstruktion der Items zurückzuführen sind, tauchten auch im Rahmen der Durchführung der Erhebung situative und auf die Anwesenheit eines Testleiters zurückzuführende Einflüsse auf die Leseleistung der SchülerInnen auf:

1) Kinder der vierten Klasse Grundschule sind bereits durch ein leistungsorientiertes Schulsystem geprägt. Demgemäß erwarteten die SchülerInnen auch in der Testsituation eine Bewertung ihrer Leistung in richtig und falsch. Zahlreiche Kinder suchten während der Testung Rückmeldung und Bestätigung ihrer Leseleistung über Blickkontakt, Fragen und genaues Beobachten der Reaktion des Testleiters. Dadurch können mögliche Auswirkungen auf die Lesezeiten, die an der Person des Testleiters und der individuellen Art der Rückmeldung liegen, nicht ausgeschlossen werden. Beispielsweise zeigten sich einige Kinder verunsichert, wenn sie erkannt hatten (beispielsweise durch ein fehlendes Nicken), dass sie ein Wort falsch gelesen hatten, und lasen folgende Wörter vorsichtiger und dadurch auch langsamer.

(2) Ein „Entziehen des Kontaktangebots“ durch keinerlei Rückmeldung oder eine automatisierte, ständig gleichbleibende Rückmeldung ist für den Umgang mit Kindern, auch in einer Testsituation, nicht zumutbar. Ein Lösungsansatz wäre mehrere Testleiter und eine größere Anzahl von Kindern für zukünftige Erhebungen heranzuziehen. Etwaige Testleitereffekte können durch unterschiedliche Testleiter moderiert werden.

(3) Auch Bewertungen der LehrerInnen können die Leseleistung beeinflussen. Beim Abholen des Kindes aus der Klasse sorgte bei manchen Kindern ein Kommentar über die schulische Leseleistung („...gehört zu den besten Lesern“. „...tut sich ein wenig schwerer“) sichtlich für Aufregung oder Entspannung. Auch wenn dies nur sehr selten auftrat, kann dieser Einfluss ohne großen Aufwand durch kurze Instruktion der LehrerInnen minimiert werden.

(4) Den Kindern wurde in der Vorbereitung auf die folgende Testsituation gesagt, dass ihre Leistungen nicht bewertet werden. Dennoch vermittelt die Art der Durchführung einen gewissen Druck, gute Leistung erbringen zu müssen/zu

wollen, alleine schon aufgrund der Instruktion, die präsentierten Wörter so schnell und so richtig wie möglich vorzulesen. Der Umgang mit Druck kann also als Drittvariable auf die Leseleistung einwirken. In der Testsituation waren unterschiedliche Verhaltensweisen der Kinder im Umgang mit der an sie gestellten Anforderung beobachtbar. Manche Kinder, besonders diejenigen, welche eine gute Sicherheit beim Lesen aufwiesen, lachten während der Testung, fanden die Pseudowörter lustig, spielten mit der Anforderung und zeigten im Verlauf der Testung keinerlei Leistungsnachlass, sondern eher einen Zugewinn. Im Gegensatz dazu zeigten sich schwache Leser oder Kinder, die sich weniger auf ihre Leseleistung zu verlassen schienen, im Umgang mit der Aufgabenstellung allgemein nervöser, vorsichtiger und verhaspelten sich häufiger. Somit kann der individuelle Umgang mit der gestellten Leistungsanforderung eine Auswirkung auf die benötigte Lesezeit haben. Hier stellt sich die Frage, ob ein Setting für zukünftige Erhebungen geschaffen werden kann, welches die Leistungsdimension minimiert. Ein spielerischer Charakter von Testungen kann den Leistungsdruck vermutlich entkräften. Bezüglich der Wahl von Methoden für Datenerhebungen an Kindern wären zukünftige Forschungen, welche speziell auf die Leistung von Kindern einwirkende Faktoren berücksichtigen und die Zumutbarkeit einer Testung gewährleisten, wünschenswert und von großem Interesse.

Eine Limitation der Studie ergibt sich aufgrund der Stichprobe, welche ausschließlich SchülerInnen im ländlichen Bereich umfasste. Des Weiteren sind die Ergebnisse vermutlich nicht übertragbar auf Kinder mit Deutsch als Zweitsprache, da nur Kinder mit Deutsch als Muttersprache untersucht wurden.

4.2 Ergebnisse

Das Ziel dieser Studie war zu untersuchen, welches Gewicht Morpheme im Prozess des lauten Lesens bei geübten Lesern haben. Dies wurde durch zwei Versuchsbedingungen getestet: durch Vorgabe von Pseudowörtern und durch silben- oder morphemorientiertes Trennen von Wörtern. Die Ergebnisse beider Bedingungen zeigen, dass morphologische Einheiten Kindern, Ende der vierten Klasse Grundschule, im Prozess des lauten Lesens zur Verfügung stehen und in diesem eine bedeutende Rolle spielen.

In beiden Bedingungen wirkte sich ein Erhalten der morphologischen Struktur positiv auf die Geschwindigkeit des lauten Lesens aus:

(1) Hauptergebnis zur Versuchsbedingung „Pseudowörter“: Die Lesezeiten zeigten, dass Pseudowörter, die aus tatsächlichen Morphemen konstruiert wurden, von den SchülerInnen signifikant schneller laut vorgelesen wurden als Pseudowörter, die davon abgeleitet wurden und keine existierenden Morpheme enthielten.

(2) Hauptergebnis zur Versuchsbedingung „Trennung“: Wörter, deren morphologische Strukturen intakt blieben, wurden von den SchülerInnen signifikant schneller laut gelesen als Wörter, deren Morpheme gestört wurden.

Gemäß dieser Ergebnisse dienen Morpheme als nützliche Verarbeitungseinheiten in der Worterkennung und werden von den Kindern als den Leseprozess unterstützende Einheiten genutzt. Dies steht in Einklang mit den Modellen der Leseentwicklung von Ehri (1998, 2005) und Klicpera et al. (2010).¹³

Die theoretische Annahme dieser beiden Modelle, dass größere Einheiten von fortgeschrittenen Lesern genutzt werden, wird insbesondere durch das oben genannte Ergebnis der Pseudowortbedingung gestützt. Ein aus Morphemen konstruiertes Pseudowort bietet den SchülerInnen die Möglichkeit, auf eine Repräsentation dieser Morpheme im Gedächtnis zuzugreifen. Ein abgeleitetes Pseudowort ohne Morphem bietet diese Möglichkeit nicht und kann nur über schrittweises Zuordnen von Graphem zu Phonem gelesen werden. Auch wenn beide Arten von Pseudowörtern in gleichem Maß artikulierbar waren, zeigte sich eine eindeutige Überlegenheit der Pseudowörter aus Morphemen im Worterkennungsprozess. Die Verfügbarkeit von morphologischen Einheiten führte zu einer signifikant höheren Lesegeschwindigkeit. Die Ergebnisse der vorliegenden Studie zeigen somit, dass Morpheme von den SchülerInnen als diese, von

¹³ In beiden Modellen wird die Annahme vertreten, dass geübte Leser Repräsentationen größerer lexikalischer Einheiten im Gedächtnis gespeichert haben und diese im Leseprozess nutzen. Kinder, welche sich in fortgeschrittenen Stadien der Leseentwicklung befinden - ab der konsolidierten Phase nach Ehri (1998, 2005) beziehungsweise ab der alphabetischen Phase mit voller Integration nach Klicpera et. al. (2010) – können beim Lesen neben Worteinträgen, auch auf andere, größere Gedächtniseinträge (Morpheme, Silben, häufige Buchstabenmuster) zurückgreifen. Insbesondere bei neuen, unbekannten Wörtern kann dies zu einer Beschleunigung des Leseprozesses führen, da diese nicht mehr ausschließlich Buchstabe für Buchstabe, sondern mithilfe des Zugriffs auf größere gespeicherte Einheiten gelesen werden können (Klicpera et al., 2010).

den beiden Entwicklungsmodellen beschriebenen, größeren Einheiten genutzt werden und den Leseprozess beschleunigen.

Die Ergebnisse der Pseudowortbedingung stehen auch nicht im Widerspruch zur Grain-Size-Theorie nach Ziegler und Goswami (2005). Kinder, die am Anfang der Leseentwicklung stehen, verlassen sich in transparenten Orthographien auf die einfache und verlässliche Zuordnung von Graphem zu Phonem, um Wörter richtig auszusprechen (Ziegler & Goswami, 2005). Kinder, Ende der Grundschule, können jedoch bereits als geübte Leser angesehen werden und im Rahmen des Zuwachs an Kompetenz somit auch größere Einheiten nutzen. Die SchülerInnen der vorliegenden Studie können beim lauten Lesen auf „Grain Sizes“ in der Größe von Morphemen zugreifen und dadurch Wörter, die aus Morphemen konstruiert wurden, schneller lesen als jene ohne morphologische Struktur. Dieses Ergebnis lässt die Annahme zu, dass sich Morpheme als verlässliche und effiziente „Grain Sizes“ im Verlauf der Leseentwicklung gebildet haben¹⁴. Die Ergebnisse der vorliegenden deutschsprachigen Studie bestätigen die Ergebnisse der Studien im italienischen Sprachraum und zeigen, dass Kinder auch in transparenten Orthographien auf größere Einheiten zurückgreifen, um den Prozess des lauten Lesens zu beschleunigen (Burani et al., 2002, 2008; Marcolini et al., 2011).

Die Verfügbarkeit von Morphemen im Leseprozess ist mit den Annahmen konnektionistischer Theorien der Worterkennung gut zu vereinen. Ein morphologisches Muster entsteht, wenn sich dasselbe orthographische oder phonologische Muster in zahlreichen, unterschiedlichen Wörtern findet und mit einer gewissen Bedeutung übereinstimmt (Plaut & Gonnerman, 2000). In der deutschen Sprache wird ein Wortstamm, der den Zugang zur Bedeutung des Wortes mit sich trägt, in all seinen möglichen Kombinationen möglichst gleich geschrieben und lässt sich somit in zahlreichen Wörtern in dieser orthographischen Form finden. Die Ergebnisse der Studie zeigen, dass dieses morphologische Prinzip der deutschen Sprache von Kindern, Ende der Grundschule, für das Lesen von unbekannten Wörtern genutzt wird. Das Vorhandensein von Morphemen in Pseudowörtern führte trotz ihrer sinnfreien und unbekannten Zusammensetzung zu kürzeren Lesezeiten.

¹⁴ Diese Annahme findet sich auch in den Studien von Burani et al. (2002, 2008).

In der Trennbedingung zeigte sich, dass ein ungestörtes Wort signifikant schneller gelesen wurde als Wörter, die ein störendes Symbol enthielten. Der Zugriff auf eine gespeicherte Gesamtwortrepräsentation beziehungsweise die Aktivierung dieses Wortmusters führte zu den signifikant kürzesten Lesezeiten. Dies steht in Einklang mit den Entwicklungsmodellen von Ehri (1998, 2005) und Klicpera et al. (2010) und auch dem DRC-Modell (Coltheart et al., 2001). Der Zugriff auf eine Repräsentation des gesamten Wortes über die direkte Route der Worterkennung stellt den effizientesten und schnellsten Weg im Leseprozess dar.

Mit dem DRC-Modell (Coltheart et al., 2001) lässt sich jedoch nicht erklären, warum Kinder in der Pseudowortbedingung Pseudowörter, die aus Morphemen zusammengesetzt wurden, schneller lasen als jene, die keine morphologische Struktur aufwiesen. Laut Zwei-Wege-Modellen existieren nur Repräsentationen ganzer Wörter als Eintrag im mentalen Lexikon. Alle Pseudowörter müssten somit über die indirekte Route der Worterkennung mittels GPK-Regeln gelesen werden, unabhängig davon, ob sie aus Morphemen konstruiert wurden oder nicht. Im DRC-Modell (Coltheart et al., 2001) ist die Aktivierung beider Routen möglich. Dadurch kann erklärt werden, dass wortähnliche Pseudowörter schneller gelesen werden als wortunähnliche. Könnten also die beobachtbaren Effekte der Pseudowortbedingung in der vorliegenden Arbeit auf Wortähnlichkeit zurückzuführen sein und nicht auf Morpheme? Diese Frage kann nicht restlos beantwortet werden. Jedoch muss zum DRC-Modell einschränkend festgestellt werden, dass dieses die Leseleistung bei einsilbigen Wörtern erklärt. Die Wortstimuli der vorliegenden Studie umfassten ausschließlich zweisilbige Wörter.¹⁵

Ein wichtiges Ergebnis der Trennbedingung der vorliegenden Studie ist, dass eine Störung der Morphemstruktur zu längeren Lesezeiten der SchülerInnen führte und somit die Worterkennung erschwerte. Eine Trennung an den Morphem-Grenzen (gleichbedeutend mit dem Erhalt der morphologischen Struktur) führte zu kürzeren Lesezeiten als eine Trennung an den Silbengrenzen (gleichbedeutend mit der Störung der morphologischen Struktur). Vereinfacht bedeutet dies auch, dass Morpheme schneller verarbeitet werden als Silben.

¹⁵ Ein Regelsystem für mehrsilbige Wörter erfordert auch entlang der indirekten Route des DRC-Modells einen Zugriff auf gespeicherte Einheiten (Rastle & Coltheart, 2000), welche Morpheme repräsentieren können.

Dieses Ergebnis ist interessant und von praktischer Relevanz, da sich der Leseunterricht in österreichischen Schulen stärker an Silben als an Morphemen orientiert. Es mag überraschen, dass sich insbesondere im Prozess des lauten Lesens, welcher eine Zuordnung von orthographischen zu phonologischen Einheiten erfordert, die morphologische Struktur der lautsprachlichen, silbischen Struktur überlegen zeigt. Zumindest in diesem Stadium der Leseentwicklung haben Kinder, Ende der Grundschule, bereits Zugriff auf morphologische Strukturen und dieser Zugriff führt zu einer effizienteren, schnelleren Verarbeitung als der Zugriff auf Silben. Mit wachsender Lesekompetenz und wachsenden Wortschatz könnten Kinder, Ende der Grundschule, zunehmend erfahren haben, dass sich morphologische Strukturen in unterschiedlichen Wörtern wiederfinden lassen, gleich bleiben und diese schließlich eine verlässlichere „Grain Size“ für die Kinder darstellen als silbische Strukturen.

Die Ergebnisse der vorliegenden Studie lassen sich am besten mit Modellen der morphologischen Verarbeitung vereinen, welche die Annahme vertreten, dass die Verarbeitung von Wörtern sowohl über Repräsentationen des gesamten Wortes als auch über jene der einzelnen Morpheme erfolgen kann (Augmented-Addressed-Morphology-Model nach Burani et al., 1984; Meta-Modell nach Schreuder & Baayen, 1995). Die Ergebnisse zeigen, dass lexikalische Worteinträge in Form von Morphemen im Gedächtnis angelegt sind. Bei unbekannten Wörtern (Pseudowortbedingung) werden diese herangezogen. Somit kann die Supralexikalische Hypothese (Giraudo & Grainger, 2000), welche besagt, dass morphologische Repräsentationen erst verfügbar sind, nachdem auf eine Repräsentation des gesamten Wortes zugegriffen wurde, nicht zur Erklärung dieses Ergebnisses herangezogen werden. In der Pseudowortbedingung kann kein Zugriff auf eine Repräsentation dieses Pseudowortes erfolgen, dennoch werden die Morpheme erkannt. Mit diesem Ergebnis im Einklang steht das Modell der obligatorischen Dekomposition (Taft & Nguyen-Hoan, 2010). Ein (Pseudo-)Wort kann auch ohne Zugriff auf eine Gesamtwortrepräsentation und ohne semantische Informationen in seine Morpheme gegliedert werden.

In der Trennbedingung führt die Möglichkeit des Zugriffs auf die Gesamtwortrepräsentation (ungestörtes Wort) zu der schnellsten Leseleistung. Zugleich wirkte eine Störung der Morpheme verzögernd auf die Leseleistung beziehungsweise ein Erhalt der morphologischen Struktur förderlich. Als Fazit kann festgehalten werden, dass

Kinder im Prozess des lauten Lesens je nach Anforderung und je nach Wortstimulus unterschiedliche Strategien nutzen: Zugriff auf Repräsentationen von Wörtern als auch Zugriff auf Repräsentationen von Morphemen.

Im Verlauf der Analyse der Tonspuren konnte eine interessante Beobachtung gemacht werden, welche einen starken Hinweis auf die Verwendung von Morphemen im Leseprozess bietet. In der Pseudowortbedingung tendierten zahlreiche Kinder bei Pseudowörtern ohne morphologische Struktur dazu, ein nicht existierendes Morphem als ein existierendes Morphem zu „erkennen“. Besonders häufig wurde das Wort „mechug“ als „mechung“ erlesen. Die nicht existierende Endung „ug“ wurde also zu der im Deutschen sehr häufig auftretenden Endung für Nomen „ung“ ergänzt. Zusätzlich zeigte sich, dass das Pseudowort „denkig“ sehr oft zu „denking“ gemacht wurde. Dies ist spannend, da Kinder Ende der Grundschule bereits Erfahrung mit der englischen Sprache haben. Das Pseudo-Wort „denkig“ wird also zu dem englischen „denking“ (thinking). Da dieser Effekt nicht bei allen Kindern im gleichen Ausmaß auftrat, stellt sich die Frage, welche Kinder ersetzen „falsche“ durch „richtige“ Morpheme in der Pseudowortbedingung. Sind es gute Leser, die über einen sicheren, automatisierten Zugriff auf Morpheme verfügen? Oder schwache Leser, welche Schwierigkeiten beim phonologischen Rekodieren aufweisen und daher auf größere, bekannte Einheiten zurückgreifen? Da im Vorfeld nicht mittels Screening zwischen guten und schwachen Lesern unterschieden wurde, bleibt diese Frage für zukünftige Forschungen zu beantworten.

In beiden Versuchsbedingungen zeigten sich keine geschlechtsspezifischen Unterschiede in den Lesezeiten der SchülerInnen. Jedoch traten klassenspezifische Unterschiede in der Pseudowortbedingung auf. Die SchülerInnen der Klasse 1 wiesen die höchste Differenz in den Lesezeiten von Pseudowörtern ohne morphologische Struktur verglichen mit Pseudowörtern aus Morphemen auf. Lassen sich diese Unterschiede in der Leseperformanz der SchülerInnen auf die erfolgte Lehrmethode zurückführen? Wurde im Unterricht das Augenmerk hauptsächlich auf größere Einheiten, wie Morpheme gelegt und weniger auf die Zuordnung von Buchstaben zu Lauten? Diese Fragen lassen sich nicht anhand der vorliegenden Daten beantworten. Zu diesem Zweck müsste im Vorfeld die angewandte Lehrmethode erhoben werden. Im folgenden Kapitel werden Überlegungen zu zukünftigen Studien näher ausgeführt.

4.3 Ausblick

Die vorliegende Studie zeigt, dass Morpheme zur effizienten Verarbeitung von Wörtern beim lauten Lesen beitragen. Aus den geschilderten Ergebnissen ergeben sich weitere Fragestellungen, welche zukünftige Studien unter Miteinbeziehung verschiedener Faktoren bedingen. Ein zusätzliches Lese-Screening im Vorfeld, welches ermöglicht zwischen guten und schwachen Lesern zu unterscheiden, könnte einen wichtigen Beitrag leisten, um praktisch relevante Implikationen für Kinder mit Lese-Rechtschreibschwäche abzuleiten. Studien im italienischen (Burani et al., 2002, 2008; Marcolini et al., 2011) als auch im dänischen (Elbro & Arnbak, 1996) Sprachraum zeigen, dass die Verfügbarkeit und Nutzung von Morphemen insbesondere schwache Leser unterstützt. Die Verarbeitung von Morphemen kann von Kindern mit LRS als eine kompensatorische Strategie genutzt werden. Zukünftige Studien können bei Kindern mit LRS erheben, welche Trennbedingung zu größeren Schwierigkeiten beziehungsweise der Erhalt welcher Struktur zu einer Erleichterung im Leseprozess führt im Unterschied zu guten Lesern. Wenn kompensatorische Strategien zur Verfügung stehen, um den Leseprozess zu beschleunigen, kann das Wissen darum in gezielten Fördermaßnahmen genutzt werden. Da Morpheme als bedeutungstragende Einheiten gelten, wäre es zusätzlich interessant, ob vorhandene Speicherung und schnellere Abrufbarkeit von Morphem-Strukturen auch zu einem besseren Verständnis des Gelesenen führen. Also ob Kinder, welche Morpheme im Leseprozess stärker nutzen, die Bedeutung des Gelesenen schneller und besser erfassen können.

In der vorliegenden Studie zeigten sich klassenspezifische Unterschiede in der Lesegeschwindigkeit je nach gestellter Anforderung. Die Bedeutsamkeit der Leseinstruktion für die Leseentwicklung wird in der Literatur mehrfach betont (z.B. Ehri, 2005; Ehri & McCormick, 1998; Klicpera et al., 2010; Ziegler & Goswami, 2005). Somit wäre es interessant für zukünftige Forschungen, im Vorfeld zu erheben, welche Leseinstruktion die Kinder in der Schule erfahren haben. Dadurch kann der Frage nachgegangen werden, wie und ob sich verschiedene Schwerpunktsetzungen im Leseunterricht (z.B. laut-, silben- oder morphemorientierter Unterricht) auf die Nutzung von orthographischen Einheiten im Leseprozess auswirken.

Diese genannten zukünftigen Forschungsfelder können einen hohen praktischen Nutzen für den Leseunterricht an Schulen darstellen. Der Leseunterricht österreichischer SchülerInnen umfasst laut Lehrplan des Bundesministeriums für Bildung und Frauen – BMBF (2012) - gegen Ende der Grundschule den Teilbereich „Sprachbetrachtung“ (BMBF, 2012, S. 129ff): Kindern soll ermöglicht werden, die Struktur der Sprache in mündlicher und schriftlicher Form zu erfassen, um dadurch mehr Sicherheit im eigenen Sprachgebrauch zu finden. Aufgaben dazu umfassen unter anderem das Bilden von Wortfamilien (z.B.: zu einem Wortstamm entsprechende Wörter sammeln) und das Zusammensetzen von Wörtern aus unterschiedlichen Wortarten und lassen somit die Bedeutung beziehungsweise Kenntnis von Morphemen für die Bewältigung dieser Aufgaben erkennen. Bredel, Noack und Plag (2011) merken jedoch kritisch an, dass das Prinzip der Stammkonstanz im Unterricht kaum systematisch erarbeitet wird. Das Konzept der Wortfamilie, welches eine Verwandtschaft von sprachlichen Ausdrücken vermitteln soll, erklärt meist nicht ausreichend, worin die „Verwandtschaft“ von Wörtern liegt, und vernachlässigt den Umstand, dass viele Schüler noch kein Konzept von lexikalischer Verwandtschaft erworben haben (Bredel et al., 2011). Die Ergebnisse der vorliegenden Studie zeigen, dass Kinder gegen Ende der Grundschule morphologische Einheiten beim Lesen nutzen und dies die Lesegeschwindigkeit erhöht beziehungsweise eine Störung dieser Strukturen zu einer Verlangsamung des Leseprozesses führt. Ob das Erkennen und Nutzen von Morphemen auf die im Unterricht erfolgte Lehrmethode zurückzuführen ist, bleibt unbeantwortet. Da jedoch die Bedeutung von Morphemen für das laute Lesen gezeigt werden konnte, sollten bestehende Leseinstruktionen evaluiert und optimiert werden.

Die vorliegende Arbeit soll und kann nicht die Frage beantworten, welches Prinzip der deutschen Sprache (silbisch, alphabetisch, morphologisch) den bedeutendsten und alles umfassenden Beitrag zu einem erfolgreichen Leseprozess leistet. In unterschiedlichen Stadien der Leseentwicklung kann jede dieser Einheiten, Buchstaben, Silben, Morpheme sowie ganze Wörter, einen wichtigen Beitrag leisten. Vielleicht führt sogar das Wissen um die eine Einheit zum Herausbilden der anderen. Es stellt sich die Frage, wie und ob sich mentale Repräsentationen von orthographischen Einheiten im Verlauf des Leseerwerbs verändern und wie sie als unterstützenden Einheiten beim Lesen dienen können je nach Lesefertigkeit und Stand der Entwicklung. Es wäre wünschenswert, dieser Frage in zukünftigen Längsschnittstudien nachzugehen.

5 Zusammenfassung

5.1 Deutsche Zusammenfassung

Morpheme gelten als die kleinsten bedeutungsunterscheidenden Einheiten der Sprache (Zwitserslood & Bölte, 2008). Das Wissen um morphologische Beziehungen ist für den Schriftspracherwerb in der deutschen Sprache von großer Bedeutung (Bredel et al., 2011; Landerl & Reitsma, 2005; Schröder-Lenzen, 2014). Das Ziel dieser Studie war zu untersuchen, welche Verwendung Morpheme im Prozess des lauten Lesens bei geübten Lesern finden. Zwei Versuchsbedingungen wurden konstruiert, um bei Kindern, Ende der Grundschule, zu erheben, ob morphologische Einheiten während dem lauten Lesen verfügbar sind und genutzt werden, um den Leseprozess zu beschleunigen. Eine Pseudowortbedingung umfasste aus Morphemen (Stamm und Suffix) zusammengesetzte Pseudowörter und dazu parallelisierte Pseudowörter ohne morphologische Struktur. In einer Trennbedingung sollten die Kinder Wörter laut vorlesen, welche mittels eines Symbols in Silben oder Morpheme segmentiert wurden. Insgesamt nahmen 38 SchülerInnen Ende der vierten Klasse zweier oberösterreichischer Grundschulen teil.

Die Ergebnisse beider Versuchsbedingungen zeigen, dass morphologische Einheiten Kindern am Ende der vierten Klasse Grundschule im Prozess des lauten Lesens zur Verfügung stehen und in diesem eine bedeutende Rolle spielen. Pseudowörter, die aus Morphemen zusammengesetzt waren, wurden von den SchülerInnen signifikant schneller laut vorgelesen als Pseudowörter, die davon abgeleitet wurden und keine Morpheme enthielten. Die Ergebnisse der Trennbedingung zeigen, dass Wörter, die eine Segmentation an den Morphem-Grenzen aufwiesen, von den SchülerInnen signifikant schneller gelesen wurden als Wörter, deren Morpheme durch eine Trennung an der Silbengrenze gestört wurden. Ein Erhalten der morphologischen Struktur führte somit zu einer höheren Lesegeschwindigkeit als ein Erhalten der silbischen Struktur.

Gemäß dieser Ergebnisse dienen Morpheme als nützliche Verarbeitungseinheiten in der Worterkennung und werden von den Kindern als den Leseprozess unterstützende Einheiten genutzt. Dies steht in Einklang mit den Entwicklungsmodellen des Lesens von Ehri (1998, 2005) und Klicpera et al. (2010), welche das Zusammenfassen von Graphem-

Phonem-Verbindungen zu größeren Einheiten mit fortschreitender Leseentwicklung annehmen. Die SchülerInnen der vorliegenden Studie können beim lauten Lesen auf mentale Repräsentationen in der Größe von Morphemen zugreifen. Dieses Ergebnis lässt die Annahme zu, dass sich Morpheme als verlässliche und effiziente „Grain Sizes“ (Ziegler & Goswami, 2005) im Verlauf der Leseentwicklung gebildet haben. Zukünftige Forschungen zu dieser Thematik sind erforderlich.

5.2 English Summary

Morphemes are deemed to be the smallest distinguishable units of any language (Zwitserslood & Bölte, 2008). The knowledge about morphological relations has been of great importance for the acquisition of reading and writing in the German language (Bredel et al., 2011; Landerl & Reitsma, 2005; Schröder-Lenzen, 2014). The objective of this study was to examine the use of morphemes in the process of reading aloud of practiced readers. Two experimental conditions were constructed to evaluate if morphological units are available to children at the end of elementary school while reading aloud, and if such are used to accelerate the reading process. A condition for pseudowords comprised pseudowords composed of morphemes (root and suffix) as well as pseudowords parallelized to the first group but without morphological structure. In a segmentation condition the children were asked to read aloud words which had been separated in syllables or morphemes by means of a symbol. In total, 38 pupils at the end of fourth grade of two elementary schools in Upper Austria took part.

The results of both experimental conditions reveal that during the process of reading aloud, morphological units are available to children at the end of elementary school and play a significant role in this process. Pseudowords composed of morphemes were read aloud considerably faster than pseudowords which were derived from the first group but included no morphemes. The results of the segmentation condition demonstrate that words which show a segmentation in morpheme clusters were read substantially faster than words in which the morphemes were deranged due to a segmentation in syllable clusters. Thus, maintaining the morphological structure led to a higher reading speed than the preservation of the syllabic structure.

According to these results, morphemes serve as useful processing units for the recognition of words and are used by children as supporting units in the reading process. This is consistent with the development models of Ehri (1998, 2005) and Klicpera et al. (2010), which hypothesize the aggregation of graphem-phonem-relations into greater units with the progressive development of reading skills. The pupils of the presented study may access mental representations in the size of morphemes during the process of reading aloud. This finding allows the hypothesis that morphemes evolve as reliable and efficient „grain sizes“ (Ziegler & Goswami, 2005) during the development of reading skills. Future research on this topic is required.

Literaturverzeichnis

- Arnbak, E. & Elbro, C. (2000). The effects of morphological awareness training on the reading and spelling skills of young dyslexics. *Scandinavian Journal of Educational Research*, 4(3), 229-251.
- Baayen, R. H., Piepenbrock, R. & van Rijn, H. (1993). The CELEX Lexical Database [CD-ROM]. Philadelphia: University of Pennsylvania, Linguistic Data Consortium.
- Baayen, H. & Schreuder, R. (1999). War and Peace: Morphemes and full forms in a noninteractive activation parallel dual-route model. *Brain and Language*, 68, 27-32.
- Bertram, R., Hyönä, J. & Laine, M. (2011). Morphology in language comprehension, production and acquisition. *Language and Cognitive Processes*, 26(4/5/6), 457-481.
- Beyersmann, E., Coltheart, M. & Castles, A. (2012). Parallel processing of whole words and morphemes in visual word recognition. *The Quarterly Journal of Experimental Psychology*, 65(9), 1798-1819.
- Bredel, U., Noack, C. & Plag, I. (2011): Morphologie lesen. Stammkonstanzschreibung und Leseverstehen bei starken und schwachen Leser/innen [online]. URL: www.phil-fak.uni-duesseldorf.de/fileadmin/Redaktion/Institute/Anglistik/Anglistisches_Institut/Dokumente/HP-data-ext/Ingo_Plag/Publications/Morphologie_lesen.pdf [14.7.2014]
- Bronk, M., Zwitserlood, P. & Bölte, J. (2013). Manipulations of word frequency reveal differences in the processing of morphologically complex and simple words in German. *Frontiers in Psychology*, 4(Article 546), 1-14.
- Bundesministerium für Bildung und Frauen (2012). Lehrplan der Volksschule, Artikel I und II, Stand BGB1. II Nr. 303/2012, September 2012 [online]. URL: www.bmbf.gv.at/schulen/unterricht/lp/lp_vs_gesamt_14055.pdf?4dzgm2 [1.7.2014]
- Burani, C., Salmaso, D. & Caramazza, A. (1984). Morphological structure and lexical access. *Visible Language XVIII* (4), 342-352.
- Burani, C., Marcolini, S. & Stella, G. (2002). How early does morpholexical reading develop in readers of a shallow orthography? *Brain and Language*, 81, 568-586.
- Burani, C., Marcolini, S., De Luca, M. & Zoccolotti, P. (2008). Morpheme-based reading aloud: Evidence from dyslexic and skilled Italian readers. *Cognition*, 108, 243-262.
- Carlisle, J. F. & Stone, C. A. (2005). Exploring the role of morphemes in word reading. *Reading Research Quarterly*, 40(4), 428-449.
- Casalis, S., Colé, P. & Sopo, D. (2004). Morphological awareness in developmental dyslexia. *Annals of Dyslexia*, 54(1), 114-138.
- Casalis, S., Dusautoir M., Colé, P. & Ducrot, S. (2009). Morphological effects in children word reading: A priming study in fourth graders. *British Journal of Developmental Psychology*, 27, 761-766.

- Coltheart, M., Perry, C. & Ziegler, J. C. (2000). The DRC Model of visual word recognition and reading aloud: An extension to German. *European Journal of Cognitive Psychology*, 12(3), 413-430.
- 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(1), 204-258.
- Ehri, L. C. (2005). Learning to read words: Theory, findings, and issues. *Scientific Studies of Reading*, 9(2), 167-188.
- Ehri, L. C. & McCormick, S. (1998). Phases of word learning: Implications for instruction with delayed and disabled readers. *Reading & Writing Quarterly*, 14, 135-163.
- Elbro, C. & Arnbak, E. (1996). The role of morpheme recognition and morphological awareness in dyslexia. *Annals of Dyslexia*, 46, 209-240.
- Giraudo, H. & Grainger, J. (2000). Effects of prime word frequency and cumulative root frequency in masked morphological priming. *Language and Cognitive Processes*, 15(4/5), 421-444.
- Goswami, U. & Ziegler, J. C. (2006). Fluency, phonology and morphology: A response to the commentaries on becoming literate in different languages. *Developmental Science*, 9, 451-453.
- Gulikers, Rattink & Piepenbrock (1993). CELEX-Manual: German Linguistic Guide [online]. URL: http://www.iub.edu/~psyling/papers/celex_eug.pdf [25. 06. 2014]
- Günther, H. (2010). Beiträge zur Didaktik der Schriftlichkeit. In H. Günther, U. Bredel & M. Becker-Mrotzek (Hrsg), *Kölner Beiträge zur Sprachdidaktik*. Duisburg: Gilles & Francke Verlag.
- Kirby, J. R., Deacon, S. H., Bowers, P. N., Izenberg, L., Wade-Wooley, L. & Parrila, R. (2012). Children's morphological awareness and reading ability. *Reading and Writing*, 25(2), 389-410.
- Klicpera, C. & Gasteiger-Klicpera, B. (1999). Lese-Rechtschreibprobleme – Einführung in den Themenschwerpunkt. *Kindheit und Entwicklung*, 8(3), 131-134.
- Klicpera, C., Schabmann, A. & Gasteiger-Klicpera, B. (2010). *Legasthenie-LRS. Modelle, Diagnose, Therapie und Förderung* (3., aktualisierte Auflage). München: Ernst Reinhardt.
- Kuo, L. & Anderson, R.C. (2006). Morphological awareness and learning to read: A cross-language perspective. *Educational Psychologist* 41(3), 161-180.
- Landerl, K. & Reitsma, P. (2005). Phonological and morphological consistency in the acquisition of vowel duration spelling in Dutch and German. *Journal of Experimental Child Psychology*, 92, 322-344.
- Marcolini, S., Traficante, D., Zoccolotti, P. & Burani, C. (2011). Word frequency modulates morpheme-based reading in poor and skilled Italian readers. *Applied Psycholinguistics*, 32, 513-532.

- 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(1), 56-115.
- Plaut, D. C. & Gonnerman, L.M. (2000). Are non-semantic morphological effects incompatible with a distributed connectionist approach to lexical processing? *Language and Cognitive Prozesses*, 15(4/5), 445-485.
- Rastle, K. & Coltheart, M. (2000). Lexical and nonlexical print-to-sound translation of disyllabic words and nonwords. *Journal of Memory and Language*, 42, 342-364.
- Rastle, K. & Davis, M. H. (2008). Morphological decomposition on the analysis of orthography. *Language and Cognitive Prozesses*, 23 (7-8), 942-971.
- Schreuder, R. & Baayen, R. H. (1995). Modeling morphological processing. In L. B. Feldman (Hrsg), *Morphological aspects of language processing*. Hillsdale, New Jersey: Lawrence Erlbaum Associates, 131-154.
- Schründer-Lenzen, A. (2013). *Schriftspracherwerb* (4., völlig überarbeitete Auflage). Wiesbaden: Springer Fachmedien.
- Steinbrink, C. & Lachmann, T. (2014). *Lese-Rechtschreib-Störung. Grundlagen, Diagnostik, Intervention*. Berlin Heidelberg: Springer-Verlag.
- Taft, M. & Nguyen-Hoan, M. (2010). A sticky stick? The locus of morphological representation in the lexicon. *Language and Cognitive Processes*, 25(2), 277-296.
- 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.

Anhang A: Pseudowort-Listen

Wort	CELEX Wortfrequenz in log/million	Pseudowort aus Morphemen	Pseudowort mit veränderten Vokalen
holen	2.08	holer	hular
können	3.625	könnung	künnang
sagen	3.31	sagig	sugog
geben	3.2	gebin	gobon
machen	3.1	machig	mechug
haben	4.00	haber	hibor
nehmen	2.91	nehmung	nihmang
denken	2.553	denkig	donkeg
brauchen	2.48	brauchin	breuchon
finden	2.896	findin	fendon

Anhang B: Wortlisten der Trennbedingung

Wörter	CELEX Wortfrequenz in log/million	Trennung in Silben	Trennung in Morpheme	Keine Trennung
fangen	1.857	fan#gen	fang#en	#fangen
merken	1.845	mer#ken	merk#en	merken#
singen	1.869	sin#gen	sing#en	#singen
lesen	2.20	le#sen	les#en	lesen#
reden	2.18	re#den	red#en	#reden
rufen	2,19	ru#fen	ruf#en	rufen#
warten	2.220	war#ten	wart#en	#warten
helfen	2.314	hel#fen	helf#en	helfen#
morgen	2.250	mor#gen	morg#en	#morgen
landen	1.532	lan#den	land#en	landen#
lenken	1.556	len#ken	lenk#en	#lenken
tanzen	1.505	tan#zen	tanz#en	tanzen#
malen	1.32	ma#len	mal#en	#malen
Hosen	1.36	Ho#sen	Hos#en	Hosen#
Regen	1.28	Re#gen	Reg#en	#Regen
sehen	3.07	se#hen	seh#en	sehen#
jeder	3.02	je#der	jed#er	#jeder
leben	2.94	le#ben	leb#en	leben#
passen	1.732	pas#sen	pass#en	#passen
retten	1.857	ret#ten	rett#en	retten#
rennen	1.838	ren#nen	renn#en	#rennen
kommen	3.227	kom#men	komm#en	kommen#
wollen	3.272	wol#len	woll#en	#wollen
müssen	3.380	müs#sen	müss#en	müssen#
Mutter	2.305	Mut#ter	Mutt#er	#Mutter
kennen	2.385	ken#nen	kenn#en	kennen#
Wasser	2.230	Was#ser	Wass#er	#Wasser
baden	1.041	ba#den	bad#en	baden#
Tafel	1.146	Ta#fel	Taf#el	#Tafel
Hasen	1.079	Ha#sen	Has#en	Hasen#

1. Block	2. Block	3. Block
fan#gen	fang#en	#fangen
merken#	mer#ken	merk#en
Sing#en	#singen	sin#gen
le#sen	les#en	lesen#
#reden	re#den	red#en
ruf#en	rufen#	ru#fen
war#ten	wart#en	#warten
helfen#	hel#fen	helf#en
morg#en	#morgen	mor#gen
lan#den	land#en	landen#
#lenken	len#ken	lenk#en
tanz#en	tanzen#	tan#zen
ma#len	mal#en	#malen
Hosen#	Ho#sen	Hos#en
Reg#en	#Regen	Re#gen
se#hen	seh#en	sehen#
#jeder	je#der	jed#er
leb#en	leben#	le#ben
pas#sen	pass#en	#passen
retten#	ret#ten	rett#en
renn#en	#rennen	ren#nen
kom#men	komm#en	kommen#
#wollen	wol#len	woll#en
müss#en	müssen#	müs#sen
Mut#ter	Mutt#er	#Mutter
kennen#	ken#nen	kenn#en
Wass#er	#Wasser	Was#ser
ba#den	bad#en	baden#
#Tafel	Ta#fel	Taf#el
Has#en	Hasen#	Ha#sen

Anhang C: Lebenslauf

Claudia Maria Colley (geb. Huber)

Geboren am 08.06.1982, Wien

AUSBILDUNG

2000	Matura, AHS Perchtoldsdorf
2000 – 2001	Universität Wien: Studium der Germanistik, der Publizistik und Kommunikationswissenschaften
Seit 2001	Universität Wien: Studium der Psychologie

FORTBILDUNGEN

2012	IVA – Instrument für Verhaltensausgleich
2012/2015	PART – Professional Assault Response Training
2015	Einblick in die Traumapädagogik

BERUFLICHE AKTIVITÄTEN

02/2015 – 08/2015	SOS-Kinderdorf, Hinterbrühl Sozialpädagogisch-therapeutische Wohngruppen Sozialpädagogin i.A.
02/2014 - 01/2015	Bildungskarenz
7-8/2011 & 12/2011 - 01/2014	Bienenhaus, Diagnose- & Therapiezentrum der SOS-Kinderdörfer, Hinterbrühl Sozialpädagogin i.A.
2001 - 2009	Österreichische Post AG, Perchtoldsdorf & Brunn am Gebirge Urlaubsvertretung/Schalterbedienstete
1998	Gemeinde Perchtoldsdorf Betreuerin bei Sommerferienspiel

PRAKTIKA UND VOLONTARIATE

10/2012-3/2013	Bienenhaus, Diagnose- & Therapiezentrum der SOS-Kinderdörfer psychologisches Praktikum
2010-2011	unterstützende Betreuung von Kindern mit tiefgreifender Entwicklungsstörung
Seit 2009	Unterstützung von Schulkindern in Gambia, Westafrika, mit Hilfe selbstorganisierter Sachspenden