



universität
wien

Diplomarbeit

Titel der Diplomarbeit

Prädiktoren des Leseverständnisses bei
Grundschulkindern der vierten Schulstufe – Erweiterung
des „simple view of reading“ Ansatzes?

Verfasserin

Stefanie Dorn

gemeinsam mit

Yvonne Huemer

Maria Klausecker

Angestrebter akademischer Grad

Magistra der Naturwissenschaften (Mag. rer. nat.)

Wien, im Oktober 2011

Studienkennzahl: 298

Studienrichtung: Psychologie

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

Danksagung

Mein Dank gilt vor allem Herrn Ao. Univ.-Prof. Mag. Dr. Alfred Schabmann für die Unterstützung und Betreuung dieser Diplomarbeit. Durch seine menschliche und unkomplizierte Art, fachkundigen Anregungen und Problemlöseansätze hat er maßgeblich zur Entstehung dieser Arbeit beigetragen.

Desweiteren möchte ich mich bei allen teilnehmenden Kindern bedanken, ohne deren Mithilfe diese Diplomarbeit nicht hätte entstehen können. Auch den Direktorinnen und Direktoren sowie den Lehrkräften der teilnehmenden Schulen danke ich für die Zusammenarbeit und Ermöglichung dieser Studie.

Der größte Dank gilt meinen Eltern und meiner Schwester Eva, bei denen ich mich herzlich für ihre Unterstützung in jederlei Hinsicht bedanken möchte. Insbesondere für die jahrelange finanzielle Unterstützung, die mir das Studium ermöglicht hat, möchte ich meinen Eltern an dieser Stelle Danke sagen.

Besonders bedanken möchte ich mich ebenfalls bei meinen Kolleginnen und Freundinnen Maria Klausecker und Yvonne Huemer für die fachlich kompetente, motivierende und von Vertrauen geprägte Zusammenarbeit im Rahmen dieses Projektes und die unzähligen schönen Erinnerungen an die gemeinsame Studienzeit. Ein großer Dank gilt auch meinem Freund Matthias Firgo für seine kompetenten Ratschläge, seine Unterstützung und sein Verständnis. Meinen Freundinnen Stefanie Drachsler und Thea Roth danke ich herzlich für ihre Ermutigungen und Lektorentätigkeit. Ein Dankeschön gebührt auch Susanne und Josef Klausecker für ihre Gastfreundschaft.

Abschließend möchte ich mich noch bei allen nicht namentlich erwähnten Freunden und Bekannten bedanken, die mich beim Schreiben der Diplomarbeit unterstützt und meine Studienzeit bereichert haben.

Inhaltsverzeichnis

Einleitung - 9 -

Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

THEORETISCHER TEIL - 11 -

1. Charakteristika der deutschen Schriftsprache - 12 -

Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

2. Definition des Leseverständnisses - 14 -

Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

3. Ebenen des Leseverständnisses - 15 -

Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

3.1. Buchstaben- und Wortebene - 16 -

Verfasst von Maria Klausecker (2011) inkl. 3.1.1 - 3.1.5.

3.1.1. Das Logogen-Modell..... - 17 -

3.1.2. Das interaktive Aktivationsmodell (Interactive-Activation-model)..... - 17 -

3.1.3. Dual-Route-Theory - 19 -

3.1.4. Dual-Route-Cascaded model..... - 20 -

3.1.5. Wortverständnis - 21 -

3.2. Satzebene - 22 -

Verfasst von Yvonne Huemer (2011)

3.3. Textebene - 24 -

Verfasst von Stefanie Dorn (2011) inkl. 3.3.1. - 3.3.3.

3.3.1. Textorientierte Ansätze - 24 -

3.3.2. Leserorientierte Ansätze..... - 26 -

3.3.2.1. *Schemageleitetes Textverstehen* - 26 -

3.3.2.2. *Vorwissen* - 27 -

3.3.2.3. *Inferenzen* - 28 -

3.3.3. Text-Leser-Interaktion: Mentale Modelle - 29 -

4. Entwicklung der Lesefähigkeit - 30 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

4.1. Der Schriftspracherwerb nach Frith und Günther..... - 31 -
Verfasst von Yvonne Huemer (2011)

4.2. Das Modell von Ehri..... - 32 -
Verfasst von Maria Klausecker (2011)

4.3. Das Kompetenzentwicklungsmodell des Lesens..... - 33 -
Verfasst von Yvonne Huemer (2011)

4.4. Simple view of reading..... - 35 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

4.5. Lexical quality hypothesis..... - 39 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

EMPIRISCHER TEIL - 43 -

5. Zielsetzung und Fragestellungen - 44 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) inkl. 5.1.
- 5.2.

5.1. Zielsetzung - 44 -

5.2. Fragestellungen..... - 45 -

6. Methodik - 48 -

6.1. Untersuchungsdurchführung..... - 48 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

6.2. Erhebungsinstrumente..... - 49 -

6.2.1. Knuspels Leseaufgaben (Knuspel-L) - 49 -
Verfasst von Yvonne Huemer (2011)

6.2.2. Ein Leseverständnistest für Erst- bis Sechstklässler (ELFE 1-6)..... - 52 -
Verfasst von Maria Klausecker (2011)

6.2.3. Peabody Picture Vocabulary Test (PPVT-III) - 54 -
Verfasst von Stefanie Dorn (2011)

6.2.4. Wortlesetest - 55 -
Verfasst von Stefanie Dorn (2011)

6.3. Untersuchungsdesign und statistische Analyse - 57 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)

7. Ergebnisdarstellung	- 60 -
Verfasst von Stefanie Dorn (2011) inkl. 7.1 - 7.2	
7.1. Deskriptive Ergebnisdarstellung.....	- 60 -
7.1.1. Stichprobe.....	- 60 -
7.1.2. Ergebnisse der Testverfahren	- 62 -
7.2. Ergebnisdarstellung der Strukturgleichungsmodelle	- 64 -
7.2.1. Modell 1	- 66 -
7.2.2. Modell 2	- 67 -
7.2.3. Modell 3	- 69 -
7.2.4. Modellvergleich	- 70 -
8. Diskussion.....	- 73 -
Verfasst von Stefanie Dorn (2011)	
8.1. Interpretation.....	- 73 -
Verfasst von Stefanie Dorn (2011)	
8.2. Entwicklungsaspekte des Leseverständnisses bei Grundschulkindern.....	- 80 -
Verfasst von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011)	
8.3. Kritik	- 82 -
Verfasst von Stefanie Dorn (2011)	
8.4. Ausblick und Relevanz.....	- 85 -
Verfasst von Stefanie Dorn (2011)	
9. Zusammenfassung.....	- 88 -
Verfasst von Stefanie Dorn (2011) inkl. 9.1. - 9.2.	
9.1. Deutsche Zusammenfassung	- 88 -
9.2. Abstract	- 89 -
10. Literaturverzeichnis.....	- 91 -
11. Abbildungsverzeichnis	- 102 -
12. Tabellenverzeichnis.....	- 103 -

13. AnhangI

Anhang A: Ansuchen zur Studienbewilligung an den LandesschulratI

Anhang B: Bewilligung der Studie durch den Landesschulrat..... III

Anhang C: ElternbriefIV

Anhang D: Testvorgabe V

Deckblatt (demographische Daten) V

Knuspel-L – Untertest Hörverstehen (Marx, 1998) VI

ELFE 1-6 – Untertest Textverständnis (Lenhard & Schneider, 2006)..... VII

adaptierte Version des PPVT-III (Dunn & Dunn, 1997) IX

Wortlesetest (Schabmann, Schmidt, Klicpera, Gasteiger-Klicpera & Klingebiel, 2009)

..... XX

Anhang E: Ergebnisse der Strukturgleichungsmodelle (Amos-Output) XXII

Modell 1XXII

Modell 2 XXIV

Modell 3 XXV

ModellvergleichXXVII

Anhang F: Curriculum vitaeXXVIII

„Wer zu lesen versteht, besitzt den Schlüssel zu großen Taten, zu unerträumten
Möglichkeiten“

–Aldous Huxley

Einleitung

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Lesen ist eine zentrale Kulturtechnik, die als elementare Voraussetzung für die Partizipation am sozialen Leben gilt. Die Fähigkeit, Texte sinngemäß zu erfassen, hat nicht nur lebenspraktische Bedeutung bei der Informationsgewinnung. Die Lesekompetenz stellt gleichzeitig eine Schlüsselqualifikation zur Wissenserweiterung dar und bedingt sowohl den schulischen, wie auch in Folge, den beruflichen Erfolg (McElvany & Schneider, 2009; Saxer, 1991). Lesekompetenz setzt sich aus mehreren abhängigen Teilfähigkeiten, die in den verschiedenen Theorien zum Leseverständnis unterschiedlich stark betont werden, zusammen und stellt somit ein komplexes Fähigkeitskonstrukt dar. Einzelne Komponenten, wie etwa Rekodieren (Fähigkeit zum Erlesen) und Dekodieren (Erkennen ganzer Wörter), Wortschatz, aber auch Hörverstehen zeichnen sich unter anderem für das Leseverständnis eines Menschen verantwortlich (Lenhard & Artelt, 2009).

Im Kontext der heutigen Mediengesellschaft, geprägt durch eine Überflutung von Informationsangeboten, kommt dem sinnentnehmenden und reflektierenden Lesen eine enorme Bedeutung in nahezu allen Lebensbereichen des modernen Menschen zu (Christmann & Groeben, 1999; Klicpera, Schabmann, Gasteiger-Klicpera, 2010). Umso problematischer erscheint vor diesem Hintergrund die Tatsache, dass vor allem im Bereich des Leseverständnisses erhebliche Defizite bei Kindern und Jugendlichen in Österreich bestehen (PISA-Studie; OECD, 2010).

Die vorliegende Diplomarbeit ist Teil eines umfassenden Projektes zur Erhebung des Leseverständnisses bei Grundschulkindern der zweiten, dritten und vierten Schulstufe. Im Zuge dieses Vorhabens sind insgesamt drei zeitgleiche Parallelarbeiten entstanden, deren Ziel in der Identifikation des Beitrags der Komponenten Worterkennen, Wortschatz und Hörverstehen für das Leseverständnis der VolksschülerInnen besteht. Ausgehend vom

„simple view of reading“-Ansatz (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000), der das Hörverstehen und die Dekodierfähigkeit als entscheidende Faktoren des Leseverständnisses postuliert, wird zusätzlich der Wortschatz, dessen ausschlaggebende Funktion im Lese(-verständnis)prozess unter anderem die „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002) betont, in der Untersuchung berücksichtigt. Die vorliegende Diplomarbeitsstudie beinhaltet die Analyse für Kinder der vierten Schulstufe, die Diplomarbeit von Klausecker (2011) beinhaltet die Analyse für Kinder der zweiten Schulstufe und die Diplomarbeit von Huemer (2011) beinhaltet die Analyse für Kinder der dritten Schulstufe.

Im theoretischen Teil der drei Arbeiten wird zunächst auf die Charakteristika der deutschen Schriftsprache und ihre Bedeutung im Leseprozess eingegangen. Daran anschließend wird das Konstrukt Leseverständnis näher erläutert. Bezug nehmend auf die einzelnen Ebenen des Leseverständnisses werden verschiedene Modelle und Theorien überblicksartig dargestellt. Das abschließende Kapitel des Theorieteils widmet sich der Entwicklung der Lesefähigkeit.

Der empirische Abschnitt umfasst die Zielsetzung und Fragestellungen der drei Parallelstudien. Auf die verwendeten Testverfahren und die Auswertung, die aus Gründen der Vergleichbarkeit der Ergebnisse ebenfalls in den Arbeiten von Huemer (2011) und Klausecker (2011) vorzufinden sind, wird im methodischen Teil eingegangen. Die Daten der SchülerInnen der vierten Klasse, die im Rahmen einer umfassenden Erhebung gewonnen wurden, werden zum einen deskriptiv statistisch dargestellt, zum anderen werden die Ergebnisse anhand von Strukturgleichungsmodellen analysiert. Schlussendlich werden die Ergebnisse im Hinblick auf den theoretischen Rahmen interpretiert, die Entwicklung des Leseverständnisses über die verschiedenen Klassenstufen hinweg dargestellt, Kritikpunkte an der durchgeführten Studie geäußert, ein Ausblick auf weitere Forschungsarbeiten gegeben und mögliche Fördermöglichkeiten des Leseverständnisses aufgezeigt.

Theoretischer Teil

1. Charakteristika der deutschen Schriftsprache

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Der Leselernprozess mit dem Ziel des Leseverständnisses setzt eine Einsicht in das jeweilige Schriftsprachsystem voraus. Ein regelmäßiges Schriftsprachsystem erleichtere dabei die tiefere Einsicht in die Phonem-Graphem-Korrespondenz (Klicpera et al., 2010). An dieser Stelle werden die Besonderheiten der deutschen Sprache kurz dargestellt.

Die deutsche Schriftsprache besteht aus Buchstaben, den fundamentalen Einheiten der alphabetischen Schrift. Sie beruht auf dem lateinischen Alphabet und ist der Gruppe der phonographischen Schriften zu zuordnen, deren Charakteristikum die Repräsentation von lautsprachlichen Elementen durch Schriftzeichen ist (Kirschhock, 2004).

Als Phonem werden die systematisch zusammengefassten Phone (beim Sprechen erzeugte Laute) bezeichnet, die kleinsten bedeutungsunterscheidenden Lautsegmente der gesprochenen Sprache. In der deutschen Buchstabenschrift werden ungefähr 40 Phoneme unterschieden. Die kleinsten bedeutungsunterscheidenden Einheiten des Schriftsystems sind Grapheme, von denen in der deutschen Schriftsprache 30 (einschließlich Umlaute und „ß“) existieren. Durch die Zuordnung von Phonemen zu Graphemen wird eine Verbindung zwischen gesprochener und geschriebener Sprache hergestellt (Scheerer-Neumann, 1997). Das asymmetrische Verhältnis verdeutlicht, dass die sogenannte Phonem-Graphem-Korrespondenz nicht eindeutig ist, da nicht alle Sprachlaute, die sich in phonologischen Merkmalen unterscheiden lassen, durch ein Buchstabenzeichen repräsentiert werden können. Verschiedene Phoneme (gekennzeichnet durch „/ /“) müssen folglich einem Graphem (gekennzeichnet durch „< >“) zugeordnet werden. Als Beispiel für diese phonetische Mehrdeutigkeit sei das Graphem <e> erwähnt, das geschlossen und kurz /e/, wie in „Religion“, aber auch als /e:/ (langes e), wie bei der ersten Silbe in „Besen“, ausgesprochen werden kann. Allerdings kann auch ein Phon durch mehrere Grapheme repräsentiert werden (graphemische Mehrdeutigkeit). Hierbei wird das /e:/ durch die Grapheme <e> (z.B. Leben), <eh> (z.B. Lehm) oder <ee> (z.B. Meer) dargestellt (Kirschhock, 2004, S. 20).

Klicpera und Gasteiger-Klicpera (1998) sehen die Regelmäßigkeit in der Phonem-Graphem-Korrespondenz als wesentliches Unterscheidungsmerkmal alphabetischer Schriften an. Im Gegensatz zur englischen oder französischen Sprache ist im Deutschen die Zuordnung von Phonemen zu Graphemen recht konsistent. Dies gilt auch für die serbokroatische oder italienische Sprache (Ziegler, Perry, Jacobs & Braun, 2001). Im Vergleich dazu können im Englischen einzelne Buchstaben und Buchstabencluster auf verschiedenen Art und Weise ausgesprochen werden beziehungsweise werden die Wörter abhängig von ihrer Aussprache unterschiedlich geschrieben (Ziegler, Stone & Jacobs, 1997). Für die deutsche Sprache gibt Naumann (1989) für 73% der Wörter eine lautgetreue Schreibweise an, was die herausragende Bedeutung des phonematischen Prinzips (Repräsentation der Phoneme in alphabetischer Schrift) widerspiegelt. Für den Leseprozess ist die „Übersetzung“ der Grapheme in Phoneme relevant. Auch hier zeichnet sich das Deutsche durch seine Regelmäßigkeit in der Graphem-Phonem-Korrespondenz aus. Es lassen sich allgemeine Regeln für die Zuordnung von Schriftzeichen in Laute (Graphem-Phonem-Korrespondenz-Regeln) ableiten, was wiederum das phonologische Rekodieren (Rekodieren von Buchstaben in lautsprachliche Form) von Wörtern erleichtert. Jedoch kann die phonetische Mehrdeutigkeit Probleme beim (lauten) Lesen bereiten, wobei die falsche Betonung von Silben schwerwiegender für den Verständnisprozess zu sein scheint. Die Aussprache von (gelesenen) Wörtern erfolgt in Silben, innerhalb derer die Laute miteinander verschmelzen. Um ein sinnentnehmendes Lesen zu gewährleisten, müssen die Grapheme je nach Stellung in der Silbe anders ausgesprochen werden (Kirschhock, 2004; Klicpera & Gasteiger-Klicpera, 1998).

Das morphematische Prinzip, das die Gleichschreibung stammverwandter Wörter regelt, stellt neben dem phonematischen das wichtigste Prinzip der deutschen Sprache dar. Als Morphem werden die kleinsten bedeutungstragenden Einheiten der Schriftsprache bezeichnet, wobei in der deutschen Sprache etwa 5000 Grundmorpheme existieren. Der große Sprachschatz im Deutschen resultiert unter anderem aus der Möglichkeit, sogenannte Stammmorpheme mit verschiedenen Affixen abzuleiten (z.B. „Gärtner-in“, „Ver-abschied-ung“) und mehrere Morpheme zu einem Wort zusammen zufassen (z.B. „Sonnen-licht“). Durch das morphematische Prinzip bleiben, trotz möglicher veränderter Aussprache, die Ableitungsformen zusammengesetzter Wörter in ihrer Schreibweise relativ konstant und selbst lange Wörter können schnell erkannt und in ihrer Bedeutung erfasst werden. Dies kann

allerdings zu Einschränkungen in der lautgetreuen Schreibweise führen (Kirschhock, 2004; Klicpera & Gasteiger-Klicpera, 1998).

2. Definition des Leseverständnisses

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Ziel der drei Parallelstudien ist, die Prozesse, die das Leseverständnis bedingen, zu untersuchen. Es stellt sich daher zu Beginn die Frage wie Leseverständnis definiert werden kann.

Das Leseverständnis stellt ein mehrdimensionales Konstrukt dar, das sich aus zahlreichen Teilfertigkeiten und Prozessen zusammensetzt. Ganz allgemein bezeichnet das Leseverständnis die Fähigkeit, Texten Informationen entnehmen zu können. Das Leseverständnis wird einerseits durch verschiedene Textmerkmale (z.B. typographische Textgestaltung und Textverständlichkeitsdimensionen) beeinflusst. Andererseits sind neben basalen Lesefertigkeiten wie Rekodieren (Erlesen von Wörtern) und Dekodieren (Erkennen von Wörtern) verschiedenste komplexe Fertigkeiten des Lesers/der Leserin notwendig, um Texte mit Verständnis lesen zu können. Die heutige Sichtweise, die durch die kognitive Psychologie und die experimentelle Leseforschung geprägt ist, fasst sinnentnehmendes Lesen als einen aktiven Prozess der Auseinandersetzung mit dem Text auf. Die im Text enthaltenen Informationen müssen weiter ausgearbeitet, miteinander verbunden und aufgrund des individuellen Vorwissens interpretiert werden (Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010; Lenhard & Artelt, 2009; Rost & Buch, 2010).

Der Leseprozess kann als Interaktion zwischen Text und LeserIn aufgefasst werden, sodass in diesem Zusammenhang textgeleitete, automatisierte „bottom-up“- und erwartungsgeleitete „top-down“-Verarbeitungsprozesse relevant sind. Von visuellen Reizen ausgehend wird sukzessiv eine kognitive Textrepräsentation aufgebaut, sodass komplexe Verständnisstrukturen entstehen können („bottom-up“-Prozess). Gleichzeitig wird das Leseverständnis durch das Interesse, die Lebenserfahrung und das spezifische Vorwissen des Lesers/der Leserin „top-down“ beeinflusst. Dieser Prozessorientierung, die ihren Fokus auf

das Zustandekommen und den Verlauf des Leseverstehens richtet und dabei die kognitiven Vorgänge berücksichtigt, steht eine Sichtweise gegenüber, die das Leseverständnis als Produkt auffasst. Das Leseverständnis als Resultat des Verstehensprozess kann durch Testverfahren erhoben werden, wobei auch seine Vorhersage durch Drittvariablen von Interesse ist (Christmann & Groeben, 1999; Rost & Buch, 2010).

Rein formell ist vom Leseverständnis das Konstrukt der Lesekompetenz zu unterscheiden. Dieses liegt der internationalen Schulleistungsvergleichsstudie PISA und der Grundschulstudie IGLU zugrunde und umfasst zusätzlich zur Verständnisfähigkeit (geschriebener) Texte auch die Fähigkeit zur Textreflektion. Es wird explizit betont, dass sich durch diese Fähigkeiten das eigene Wissen progressiv erweitern kann und eine Teilnahme am gesellschaftlichen Leben ermöglicht wird. Die Lesekompetenz kann somit als eine umfassende Erweiterung des Begriffs „Leseverständnis“ aufgefasst werden (OECD, 2010; Rost & Buch, 2010).

3. Ebenen des Leseverständnisses

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Das Leseverständnis, das im Rahmen dieser Diplomarbeiten untersucht wird, resultiert aus einem komplexen Wechselspiel zwischen verschiedenen Teilkomponenten (Klicpera et al., 2010). Jene Bausteine, die sich für das verstehende Lesen verantwortlich zeichnen, sollen daher an dieser Stelle dargestellt werden.

Für das Verständnis von Texten und für die Bedeutungsentnahme sind verschiedene Verarbeitungsschritte notwendig, die sich auf unterschiedlichen Komplexitätsebenen (Buchstaben- und Wortebene, Satzebene und Textebene) vollziehen und dabei verschiedenste Anforderungen an die Kompetenzen des Lesers/der Leserin stellen (Christmann & Groeben, 1999; Klicpera et al., 2010; Lenhard & Artelt, 2009).

3.1. Buchstaben- und Wortebene

Dieser Abschnitt wurde verfasst von Maria Klausecker (2011).

Um das Ziel des Lesens, das Textverständnis, zu erreichen, muss zuerst die Bedeutung einzelner Wörter erfasst werden. Basale Wahrnehmungsprozesse gehen dabei dem Erlesen eines Wortes voraus. In einem ersten Schritt müssen visuelle Reize wahrgenommen und verarbeitet werden. Innerhalb eines Lesevorgangs werden einzelne Wörter mit den Augen fixiert und dann wird mit einer kurzen Augenbewegung (Sakkade) zum nächsten Wort übergegangen. Das Lesen ist folglich für den geübten Leser ein Wort für Wort Lesen (Klicpera & Gasteiger-Klicpera, 1998). Es konnte gezeigt werden, dass anhand der Augenbewegungen der Leseprozess gut untersucht werden kann (Rayner, 1998).

Der sogenannte Wortüberlegenheitseffekt beschreibt ein Phänomen, dass Buchstaben innerhalb von bekannten Wörtern schneller wahrnehmbar sind als Buchstaben in unsinnigen Buchstabenfolgen oder auch einzelne Buchstaben. Dieser Effekt spricht gegen die augenscheinliche Vermutung, dass die Worterkennung über das Erlesen einzelner Buchstaben erfolgt. Auch sogenannte Pseudowörter, aussprechbare sinnlose Buchstabenaneinanderreihungen, weisen einen Vorteil gegenüber zufälligen Buchstabenfolgen auf und können in etwa ebenso gut wahrgenommen werden wie echte Wörter. Es wird geschlossen, dass der Wortüberlegenheitseffekt mit der Regelmäßigkeit der Schriftsprache und nicht mit der Vertrautheit bestimmter Wörter zusammenhängt. Heute wird davon ausgegangen, dass der Wortüberlegenheitseffekt durch die Verfügbarkeit eines schnell zugänglichen und relativ stabilen Kodes, der das Behalten der vorgegebenen Information erleichtert, entsteht (Klicpera & Gasteiger-Klicpera, 1998).

Die Ebene der Worterkennung wird auch durch den Kontext beeinflusst (Oakhill & Garnham, 1988): Der Kontext kann den Leser bei der Korrektur oder Vermeidung von Lesefehlern unterstützen. Die Worterkennung wird zusätzlich durch die Voraktivierung ähnlicher Bedeutungsinhalte beschleunigt. Das Lesen auf Wortniveau ist durch den Kontext, wie etwa den Satz, in dem das Wort eingebettet ist, beeinflusst.

Aus diesen Befunden leiteten sich einflussreiche Theorien und Modelle ab. Hier wird auf vier Theorien näher eingegangen, die die Forschung stark beeinflusst haben: Das „Logogen-Modell“ von Morton (1979), das „interaktive Aktivationsmodell“ nach McClelland und

Rumelhart (1981), die „Dual-Route-Theory“ (Coltheart, 1978) und das „Dual-Route-Cascaded model“ (DRC-Modell, Coltheart, Rastle, Perry, Langdon, & Ziegler, 2001).

3.1.1. Das Logogen-Modell

Ausgehend von seinem Modell über den Worterkennungsvorgang beim Lesen (1969) entwickelte Morton das Logogen-Modell (1979), das die ablaufenden Lese- und Schreibprozesse beim geübten Leser/bei der geübten Leserin darstellt und das Konzept eines inneren Wortspeichers beziehungsweise Lexikons einführt (zitiert nach Klicpera & Gasteiger-Klicpera, 1995, S. 97ff.).

Das Modell, das durch die Informationsverarbeitungstheorie geprägt ist, weist die typische Basisstruktur „Input → Verarbeitung → Output“ auf (Graf, 1994). Morton (1979) unterscheidet zwischen dem Eingangs-Logogen-System, den akustischen/visuellen Worterkennungseinheiten, dem Ausgangs-Logogen-System, den akustischen/visuellen Worterzeugungseinheiten und dem kognitiven System. Als Logogene werden mentale Einheiten bezeichnet, in denen alle relevanten Informationen über die Schreibweise, Aussprache, Funktion im Satzkontext und Bedeutung des Wortes gespeichert sind. Entscheidend für die Verarbeitung von Schrift ist das kognitive System, das den Input (Schrift) über verschiedene Prozesse in den Output (Erkennen von Wörtern und Sätzen beziehungsweise Aussprache) überführt (Klicpera et al., 2010).

Informationen sind in diesem Modell die in den Wörtern enthaltenen Buchstaben, die im Verlauf der visuellen Analyse im kognitiven System eintreffen. Dadurch wird das entsprechende Logogen aktiviert und sobald die Aktivierung einen bestimmten Schwellenwert überschreitet, wird das Wort erkannt.

3.1.2. Das interaktive Aktivationsmodell (*Interactive-Activation-model*)

McClelland und Rumelhart (1981) gehen davon aus, dass nicht nur Buchstaben und Wörter im Gedächtnis gespeichert sind, sondern auch grafische Elemente, also Punkte und Striche, aus denen sich die einzelnen Buchstaben zusammensetzen. Diese gesamte Information ist im Gedächtnis als Netzwerk gespeichert. Die Aufnahme der grafischen Elemente startet bereits den Worterkennungsprozess und löst typische Aktivierungsmuster im neuronalen Netz aus: Die Buchstaben, die den Merkmalen entsprechen, werden durch exzitatorische Verknüpfung

aktiviert, andere über inhibitorische Verbindungen gehemmt. Dieses System lässt sich auch auf die Buchstaben- und Wortebene übertragen. Bei Erreichen einer bestimmten Aktivitätsschwelle, wird der/das entsprechende Buchstabe/Wort aktiviert.

Es wird neben dem postulierten „bottom-up“-Prozess (ausgehend von der untersten Ebene/Sinnesreiz bis zur höchsten Ebene/Wort) zusätzlich angenommen, dass die Informationsverarbeitung auch „top-down“ (lexikalisch) geleitet wird. Vermutungen, die aus dem Vorwissen über Wörter resultieren, beeinflussen die Wahrnehmung der verbleibenden Buchstaben („top-down“-Prozess). Daraus lässt sich schlussfolgern, dass die Schreibweise von bereits bekannten Wörtern (der orthographische Kode) in einem inneren Lexikon gespeichert sein muss, auf das direkt zugegriffen werden kann. Das Zusammenspiel von „top-down“- und „bottom-up“-Prozessen beim Lesen bewirkt die Effizienz der Worterkennung (Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010).

Negativ an diesem Modell anzumerken ist die Tatsache, dass wie auch im Logogen-Modell (Morton, 1979) vorausgesetzt wird, dass jedes Wort als Einheit in einem mentalen Lexikon gespeichert ist. Die Identifikation neuer oder zusammengesetzter Wörter oder auch das Erkennen verschiedener Flexionsformen kann durch dieses Modell nicht erklärt werden. Die Annahme, dass neben dem direkten lexikalischen Weg ein weiterer indirekter Weg bestehen muss, ist daher leicht nachzuvollziehen (vgl. Coltheart, 1978).

Eine Weiterentwicklung des Ansatzes von McClelland und Rumelhart (1981) stellt das „parallel-distributed-processing (PDP) connectionist model“ von Seidenberg und McClelland (1989) dar. Sie gehen davon aus, dass das Wissen über Wörter im Laufe der Erfahrung mit der Schriftsprache in Form von neuronalen Netzwerken aufgebaut wird, auf die schnell und unbewusst zugegriffen werden kann. Es existiert kein einheitliches mentales Lexikon, sondern die Netzwerke enthalten Einheiten, welche die orthografischen, phonologischen und semantischen Eigenschaften von Wörtern repräsentieren. Über sogenannte verborgene Einheiten („hidden units“), die vom Netzwerk durch Erfahrung/Übung gebildet werden, sind die einzelnen Komponenten miteinander verbunden. Das Wissen über Wörter ist durch die probabilistischen Verbindungen bedingt, die zu einer Aktivierung der unterschiedlichen Repräsentationseinheiten führen (Kirschhock, 2004; Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010; Seidenberg & McClelland, 1989).

Durch einen visuellen Input (Wort) werden verschiedene orthografische Einheiten (Eingang) im Netzwerk aktiviert, die mit den verdeckten Einheiten verbunden sind. Das Aktivierungsmuster breitet sich auf die phonologischen Einheiten (Ausgang) aus und in mehreren Zyklen wird durch rückkoppelnde Schritte schließlich eine Gesamtaktivierung erreicht, die die Identifikation der Buchstabenfolgen eines Wortes ermöglicht. In der Ausgangseinheit ist dann das zu lesende Wort in seiner Lautsprache repräsentiert.

Am Computer können die postulierten Prozesse überprüft werden, wobei ein Programm die Rolle des Menschen beim visuellen Worterkennen und lauten Lesen simuliert. Je vergleichbarer die erzielten Leistungen des Programms mit denen von Menschen sind, umso eher kann davon ausgegangen werden, dass die angenommenen Modellannahmen beim Menschen auf korrespondierende Weise ablaufen. Die experimentelle Überprüfung, bei der ein Netzwerk mit 400 orthografischen Einheiten, 200 „hidden units“ und 460 phonologischen Einheiten am Computer simuliert wurde, ermittelte für 90% der Pseudowörter eine korrekte Aussprache. Dieses Ergebnis lässt darauf schließen, dass sich die angenommenen Einheiten des PDP-Modells bewähren, um den Vorgang des lauten Lesens zu simulieren. Jedoch fehlen für die Annahme, dass auch Wortbedeutungen und Kontexteinflüsse in einem Netzwerk aktiviert werden können, bislang die empirischen Belege (Lenhard & Lenhard, 2009c).

Computerprogramme von konnektionistischen Modellen treffen Entscheidungen aufgrund „gelernter“ Muster. Allerdings weisen schwache Leser oder Leseanfänger solche Muster nicht auf, sodass der Beginn des Schriftspracherwerbs nur unzureichend mit solchen Modellen beschrieben werden kann (Mannhaupt, 2001).

3.1.3. Dual-Route-Theory

Die auf Coltheart (1978) zurückgehende Theorie der visuellen Worterkennung nimmt zwei verschiedene Verarbeitungsrouten (auch „dual-route-theory“ genannt) im Leseprozess an. Der direkte, lexikalische Weg besteht in der orthografischen Kodierung des Schriftbildes, die den entsprechenden Eintrag im mentalen Lexikon aktiviert, sodass die Aussprache unter Zugriff auf die phonologisch-artikulatorische Information sofort möglich wird. Bei unbekannten Wörtern oder Pseudowörtern muss die indirekte, phonologische Route gewählt werden. Da hier kein Eintrag im Lexikon vorhanden ist, wird das Wort seriell durch phonologisches Rekodieren (Graphem-Phonem-Zuordnung) erschlossen, was durch die Regelmäßigkeit der deutschen Schriftsprache möglich wird. Vermutlich bauen LeseanfängerInnen über den

indirekten Weg ihr mentales Lexikon auf und je umfangreicher dieses ist, desto leichter und häufiger wird der direkte Abruf möglich (Coltheart & Rastle, 1994; Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al. 2010; Stock, 2009).

Beide Wege sind für den Leseprozess essentiell, wie neuropsychologische Befunde von PatientInnen mit erworbener Leseschwäche verdeutlichen. Im Fall der Oberflächendyslexie ist der direkte Weg durch hirnorganische Schädigungen beeinträchtigt und Wörter müssen Graphem für Graphem rekodiert werden. Dadurch unterbleibt eine notwendige Automatisierung des Lesevorgangs und das Erlesen von unregelmäßigen und homophonen Wörtern (Wörter mit gleicher Phonologie und verschiedener Bedeutung) bereitet Schwierigkeiten. Ist hingegen der indirekte Weg beeinträchtigt, kommt es zu Defiziten beim phonologischen Rekodieren, was als phonologische Dyslexie bezeichnet wird. Vor allem unbekannte Wörter (z.B. Pseudowörter) können schwer gelesen werden, da kein Eintrag im mentalen Lexikon bereitsteht (Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010; Lenhard & Artelt, 2009).

Scheerer-Neumann (1997) weist in Untersuchungen darauf hin, dass die beiden Wege der „dual-route-theory“ nicht unabhängig voneinander verlaufen. Es wird davon ausgegangen, dass beim Lesen neuer Wörter lexikalisches Wissen mitbenutzt wird.

3.1.4. Dual-Route-Cascaded model

Das „Dual-Route-Cascaded model“ (DRC-Modell, Coltheart et al., 2001) hat sich in Computersimulationen sehr gut bewährt. Laut diesem Modell läuft der Leseprozess in den folgenden Schritten ab. Zuerst werden die visuellen Eigenschaften eines Wortes analysiert. Im nächsten Schritt werden die lexikalische und die nicht lexikalische Route des Modells gestartet. Die nicht lexikalische Route beginnt die Umwandlung von Graphemen in Phoneme. Über den lexikalischen Weg lösen die Buchstabenfolgen die Aktivierung des gesamten Wortes im orthographischen Lexikon aus. Jene Route die als erstes zu einem Ergebnis kommt „gewinnt“ und ermöglicht dann die korrekte Aussprache des Wortes. In einer Erweiterung dieses Modells wird auf lexikalischer Ebene noch ein „semantisches System“ eingeführt, welches die Wortbedeutungen beinhaltet. Es steht mit dem orthographischen und phonologischen Lexikon in Verbindung (Klicpera et al., 2010).

Das DRC-Modell wurde auf die deutsche Sprache übertragen und konnte seine Gültigkeit in der deutschen Sprache auch belegen (Ziegler, Perry & Coltheart, 2000). In einem experimentellen Versuchsdesign wurde die Verarbeitung des non-lexikalischen Weges im DRC-Modell allerdings kritisiert. Bei einem Vergleich der Leseleistungsmuster der Computer-Simulation mit der von 24 Versuchspersonen der Universität Waterloo ergaben sich qualitative Unterschiede. Es wird eine Überarbeitung der nicht-lexikalischen Route im DRC-Modell verlangt (Besner & Roberts, 2003).

Eine Verbesserung des DRC-Modells stellt das „CDP+ Model of Reading Aloud“ (Perry, Ziegler & Zorzi, 2007) dar. Einen wesentlichen Vorteil des „CDP+ Model of Reading Aloud“ gegenüber dem DRC-Modell stellt die Tatsache dar, dass das CDP+ Modell fähig ist, neue Wörter zu lernen und gestaffelte Konsistenzeffekte produzieren kann.

3.1.5. Wortverständnis

In den bereits beschriebenen Modellen wird das Wortverstehen als Zugriff auf gespeicherte Informationen über ein Wort im semantischen Lexikon beziehungsweise in Netzwerken verstanden. Klicpera et al. (2010) gehen davon aus, dass das Bedeutungswissen über Wörter mit anderen Fähigkeiten (z.B. Vorwissen, Wortschatz) des Leseverständnisses und dem Kontext interagiert. Obwohl der Einfluss des Kontexts auf Wortebene umstritten ist (für einen Überblick siehe Rayner & Pollatsek, 1989), scheint die Annahme gerechtfertigt, dass Kontextinformationen für das Verständnis durchaus eine wichtige Funktion inne haben, den visuellen Worterkennungsprozess per se jedoch nicht beeinflussen (Klicpera & Gasteiger-Klicpera, 1998).

Bedeutend ist in diesem Zusammenhang die Kodierung morphologischer Strukturen von Wörtern, was als dritte Möglichkeit der Wortidentifikation angesehen werden kann (für den deutschsprachigen Raum siehe „Hypothesentest- und Redundanzausnützungsmodell“ von Grisse, 1996). Hierbei wird zunächst der visuelle Input in morphologische Einheiten zergliedert und nach der Identifikation des Stamm-Morphems erfolgt die Analyse der Affixe. Dementsprechend werden auch zusammengesetzte Wörter erkannt, wobei für jedes Stamm-Morphem ein Eintrag im mentalen Lexikon angenommen wird. Die Ergebnisse von Zwitserlood (1996) deuten darauf hin, dass die morphematische Organisation des Lexikons insbesondere für die Bedeutungsermittlung von Wörtern ausschlaggebend ist. Das Wortverständnis wird durch das Erkennen und Reflektieren der Zusammensetzung von

Wortstämmen zu neuen Wörtern erheblich erleichtert (Christmann & Groeben, 1999; Klicpera & Gasteiger-Klicpera, 1998).

3.2. Satzebene

Dieser Abschnitt wurde verfasst von Yvonne Huemer (2011).

Um die Bedeutung eines Satzes zu verstehen, müssen die identifizierten Wörter und deren Bedeutungen zueinander in Beziehung gesetzt und in ein Gesamtgefüge eingegliedert werden. Dies setzt eine Analyse der semantischen und syntaktischen Relationen der einzelnen Satzteile voraus (Christmann & Groeben, 1999).

Der semantische Bedeutungsgehalt wird durch die Extraktion von Propositionen gewonnen. Sätze sind nach dem Prinzip der Prädikat-Argument-Struktur (Propositionen) aufgebaut, wobei das Prädikat die semantischen Relationen festlegt und an der Satzoberfläche als Verb, Adjektiv oder Adverb in Erscheinung treten kann. Der Satz „Die Katze beißt den Hund“ enthält beispielsweise die Propositionen BEISSEN, KATZE, HUND (Christmann & Groeben, 1999, S. 153). Die semantische Analyse kann alleine keine eindeutige Bedeutungszuordnung gewährleisten (so ist für das Beispiel aufgrund der Propositionen nicht ersichtlich, ob die Katze den Hund beißt oder umgekehrt) und muss durch die syntaktische Analyse ergänzt werden. Hierbei werden Wörtern und Wortgruppen syntaktische Funktionen wie Subjekt, Prädikat und Objekt zugewiesen und die Sätze im einfachsten Fall anhand der kanonischen Sentoid-Strategie (Fodor, Bever & Garrett, 1974; zitiert nach Christmann & Groeben, 1999, S. 154) erschlossen. Für semantisch eindeutige Sätze (Sentoid) ergibt sich der Bedeutungsgehalt durch die Abfolge der Inhaltswörter (Katze = Subjekt, beißen = Prädikat, Hund = Objekt). Bei komplexeren Sätzen und Passivsätzen („Der Hund wird von der Katze gebissen“) ist dieses Vorgehen jedoch nicht erfolgreich, da die Reihenfolge der Wörter (Oberflächenstruktur) nicht der tatsächlichen grammatikalischen Bedeutung (Tiefenstruktur) entspricht und weitere syntaktische Informationen zur Bedeutungsanalyse genutzt werden müssen (Christmann & Groeben, 1999).

Bei der als „parsing“ bezeichneten Tiefenstrukturanalyse wird jedem gelesenen Wort automatisch eine bestimmte Position im Satz zugewiesen. Diese Platzierung ergibt sich aus der Annahme der wahrscheinlichsten Satzkonstruktion, die bei Hinzukommen widersprüchlicher Informationen revidiert wird. Insbesondere bei syntaktisch ambigen

Sätzen kommt es durch das inkrementelle Vorgehen zu einer falschen syntaktischen Strukturierung, die im Verlauf der Verarbeitung korrigiert werden muss (sog. „garden-path“-Effekt). Als neuronales Korrelat wurde dabei von Osterhout, Holcomb und Swinney (1994) die P600-Komponente identifiziert. Dieses ereigniskorrelierte Potential konnte bei Probanden dann gemessen werden, wenn die grammatikalische Struktur schwer zu verstehen beziehungsweise verletzt war. Als zentrale Gliederungskomponente für Sätze gelten Nominalphrasen, die Nomen, Artikel und Adjektive umfassen, sowie Verbalphrasen, die ein Verb und weitere Nominalphrasen enthalten (Klicpera et al, 2010; Lenhard & Lenhard, 2009b). Die syntaktische Analyse beruht dem „garden-path“-Modell zufolge auf zwei wesentlichen Grundprinzipien. Zum einen wird die zu rekonstruierenden Satzstruktur so gebildet, dass sie möglichst wenige Verzweigungen aufweist („minimal attachment“), zum anderen wird das gerade gelesene Wort nach Möglichkeit in die zuletzt aktive Phrase eingebaut („late closure“) (Klicpera & Gasteiger-Klicpera, 1998). Dies führt laut Christmann und Groeben (1999) dazu, dass der Verarbeitungsaufwand beim Lesen minimiert und die Effizienz maximiert wird.

Bezüglich der Priorität der Syntax bei der Satzsegmentierung oder der Interaktion von Semantik und Syntax besteht Dissens in der Literatur, wobei sich zwei gegensätzliche Positionen unterscheiden lassen. Die interaktionistische Syntaxtheorie geht davon aus, dass die Syntaxanalyse vom semantischen und pragmatischen Kontext sowie vom Weltwissen abhängt, und syntaktische und semantische Teilprozesse parallel ablaufen. Die autonome Syntaxtheorie postuliert die Unabhängigkeit der beiden Prozesse und nimmt an, dass die syntaktische Verarbeitung der semantischen zeitlich voran geht. Zahlreiche Befunde weisen darauf hin, dass keine der beiden Theorien absolute Gültigkeit für sich beanspruchen kann. Bei syntaktisch ambigen Sätzen oder bei geringem kontextuellen Bezug wird beim geübten Leser/bei der geübten Leserin gemäß der autonomen Syntaxtheorie die grammatikalische Struktur des Satzes isoliert verarbeitet. Liegen eindeutige Kontextbezüge vor beziehungsweise entsteht während des Parsens kein Widerspruch, so beschreibt die interaktionistische Syntaxtheorie den Leseprozess besser. Insgesamt scheint die Annahme berechtigt, dass bei der Satzanalyse semantische Sinnesstrukturen mit Hilfe der Syntax aufgebaut werden, wobei umso mehr auf die Syntax zurückgegriffen wird, je komplexer die Sätze sind. Wurde die Satzbedeutung erfasst, wird die syntaktische Information vergessen. Vermutlich sind beide Ansätze beim Leseprozess wichtig und welche Vorgehensweise

verwendet wird, ist vom Leseverständnis und der Textschwierigkeit abhängig (Christmann & Groeben, 1999; Lenhard & Artelt, 2009).

Bei Kindern, die sich in der Sprachentwicklungsphase befinden, scheint eine relative Autonomie der syntaktischen Analyse vor zu herrschen. Ein wichtiger Faktor bezüglich des Leseverständnisses von Kindern ist Klicpera und Gasteiger-Klicpera (1998) zufolge in der Entwicklung der Fähigkeit zur Satzanalyse beziehungsweise der grammatikalischen Kompetenz zu sehen.

3.3. Textebene

Dieser Abschnitt wurde verfasst von Stefanie Dorn (2011).

Das Textverständnis ist durch zwei strategische Prozesse bedingt, die auch die Forschungsansätze auf diesem Gebiet bestimmen. So geht der textseitig orientierte Ansatz davon aus, dass die Informationsaufnahme von den elementaren Aussagen bis zur Rekonstruktion des Themas fortschreitet. Dementsprechend wird versucht, Textstrukturen objektiv zu beschreiben und die Auswirkungen von Texteigenschaften auf das Lesen, Verstehen und Behalten zu ermitteln. Ausgehend vom Verständnis, fokussiert der leserseitig orientierte Ansatz den Einfluss von Vorwissen, Weltwissen, Erwartungen und Zielsetzung auf die Textverarbeitung (Christmann & Groeben, 1999; Klicpera & Gasteiger-Klicpera, 1998).

3.3.1. Textorientierte Ansätze

Auf lokaler Ebene wird die Textbasis durch die Integration von Informationen, die in kürzeren Abschnitten enthalten sind, aufgebaut. Zunächst wird durch die Aktivierung eines Netzwerkes, in dem Konzepte für das entsprechende Wort mit Konzepten anderer Wörter verbunden sind, die Wortbedeutung herausgearbeitet. Die Aktivierung aller entsprechend verfügbaren Konzepte bei mehrdeutigen Wörtern wird schrittweise durch den Kontext reduziert, wobei Pickering und Traxler (1998) davon ausgehen, dass zunächst das Konzept mit der wahrscheinlichsten Wortbedeutung aktiviert wird. Dieser Informationsreduktionsprozess kann nach einzelnen Sätzen einsetzen oder auch erst nach Kenntnis längerer Textpassagen (Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010).

Das Modell der Propositionsanalyse (Kintsch, 1974) ist der bekannteste Ansatz zur Erklärung der Integration von Einzelinformationen. Durch das Verständnis einzelner Wörter bedingt,

werden elementare Texteinheiten oder Aussagen (Propositionen) gebildet, von denen mehrere in einem Satz vorkommen können. Propositionen sind in Form einer Prädikat-Argument-Struktur aufgebaut, wobei das Prädikat (Verb, Adjektiv) die zentrale Organisationseinheit darstellt. Ausgehend vom Prädikat werden weitere Propositionen aufgebaut (hinzufügen des Subjekts, Objekts, etc.) und schließlich kann der Satz aufgrund seiner Propositionen in eine Makroproposition integriert werden. Texte werden als Liste von Propositionen notiert, was als Textbasis bezeichnet wird (Christmann & Groeben, 1999; Klicpera & Gasteiger-Klicpera, 1998; Lenhard & Lenhard, 2009d).

Um Zusammenhänge zwischen verschiedenen Propositionen, Sätzen oder Textstellen herzustellen, werden in der Sprache sogenannte Kohäsionsmittel verwendet. Diese umfassen rückbezügliche Ausdrücke (Anaphern), wie Demonstrativpronomina, Possessivpronomina oder rückbezügliche Adverbien des Ortes und der Zeit. Ein Überblick über die verschiedenen Kohäsionsmittel und deren Auswirkungen im Verarbeitungsprozess ist bei Christmann und Groeben (1999, S. 158) zu finden.

Die Herstellung semantischer Relationen zwischen Sätzen bzw. Propositionen wird als Bildung der lokalen Kohärenz bezeichnet (vgl. van Dijk & Kintsch, 1983). Demzufolge werden Propositionen miteinander verknüpft, wenn sie die gleichen Prämissen enthalten oder wenn eine Proposition in eine andere integriert ist. Von dieser Basis ausgehend wird eine hierarchische Textstruktur erstellt. Kintsch (1974) vermutet, dass denjenigen Propositionen, die eine logische Voraussetzung für weitere Aussagen bilden, eine zentrale Organisationsfunktion, im Sinne einer hierarchischen Textnachbildung, zukommt. Dies entspricht einer Differenzierung zwischen zentralen Inhalten und Detailinformationen. (Christmann & Groeben, 1999; Klicpera & Gasteiger-Klicpera, 1998).

Christmann und Groeben (1999) würdigen den erfolgreichen Beitrag der Propositionsmodelle für die Grundlagenforschung, weisen aber gleichzeitig auf die zahlreichen Schwachstellen hin (Überblick siehe Christmann & Groeben, 1999, S. 164). In ihrem Modell der zyklischen Verarbeitung versuchen Kintsch und van Dijk (1978) die leserseitigen Wissensbestände und Inferenzen zu berücksichtigen. Allerdings ist auch diese Weiterentwicklung des Propositionsmodells (Kintsch, 1974) nicht in der Lage, den Verständnisprozess bei längeren und komplexeren Texten adäquat abzubilden. Dies soll mittels Makrostrukturmodellen realisiert werden, die eine globale Kohärenzherstellung auf höheren Abstraktionsebenen

darzustellen versuchen. Beim Erschließen längerer Texte werden die Informationen auf das Wesentliche verdichtet, sodass eine Makrostruktur entsteht. Bei der Bildung von Makrostrukturen nehmen van Dijk und Kintsch (1983) eine konstruktive Interaktion zwischen Text und Vor- und Weltwissen des Lesers/der Leserin an, womit auch bei den textorientierten Ansätzen schon die Wichtigkeit der leserseitigen Aspekte betont wird (Christmann & Groeben, 1999; Lenhard & Artelt, 2009). Die Befunde der kognitionspsychologischen Forschung, beispielsweise zu kurz- und langfristigen Erinnerungsleistungen (Kintsch & van Dijk, 1978) und Priming-Effekten (Guindon & Kintsch, 1984), belegen die Relevanz der Makrostrukturbildung als Teil des Verständnisprozesses. Allerdings stellen die Makroregeln und die ablaufenden kognitiven Vorgänge ein schwer operationalisierbares Konstrukt dar, sodass diese in der Forschung bisher wenig beachtet wurden (Christmann & Groeben, 1999).

3.3.2. Leserorientierte Ansätze

3.3.2.1. Schemageleitetes Textverstehen

Der auf Bartlett (1932) zurückgehende Begriff des Schemas repräsentiert ein abstraktes Gefüge von hierarchisch geordnetem, allgemeinem Weltwissen über typische Zusammenhänge realer Sachverhalte/Situationen. Diese ordnenden Konzepte werden durch die eng mit dem Schema verbundenen Informationen aktiviert. Sie steuern die Aufmerksamkeit und erleichtern die Integration und Interpretation neuer Informationen (detaillierte Ausführung siehe Christmann & Groeben, 1999, S. 167). Insbesondere auf der Makroebene erweist sich die Verwendung von Schemata als besonders hilfreich, da ein Modell der Situation durch die Informationen im gelesenen Text mental generiert wird und den Verständnisprozess „top-down“ steuert (Klicpera & Gasteiger-Klicpera, 1998).

Für das Textverständnis relevant ist in diesem Zusammenhang die Unterteilung des Schema-Begriffs in Skripts und Geschichtengrammatiken („story grammars“). Ein Skript umfasst das spezifische Wissen über typische Handlungsabläufe in stereotypen Situationen (z.B. Restaurantbesuch, Einkauf, etc.) und durch seine Aktivierung kann während des Lesens fehlende Information, unter Rückgriff auf das entsprechende Skript, generiert werden (vgl. Abbott, Black & Smith, 1985). Die Organisation von Textelementen in Form von Kategorien über einzelne Elemente und Aufbaueregeln von Erzähltexten wird als Geschichtengrammatik bezeichnet, die dem Leser/der Leserin ein allgemeines (mentales) Gliederungsschema für den Text zur Verfügung stellt. Geschichten, die nach dem Grundprinzip des Handlungsschemas

Ausgangssituation – Motiv – Ziel – Versuch, das Ziel zu erreichen – Ergebnis aufgebaut sind, werden auch von kleinen Kindern schon recht gut verstanden, dem Verlauf kann leichter gefolgt und die Inhalte besser erinnert werden (Christmann & Groeben, 1999; Klicpera & Gasteiger-Klicpera, 1998). Sachtexte sind ebenfalls nach bestimmten rhetorischen Strukturen aufgebaut. Sind Strukturen für Sachtexte bekannt, dann gelingt es, Informationen besser wiederzuerkennen, zu organisieren und abzurufen. Somit kommt der Textstruktur die Funktion einer Gedächtnishilfe zu (Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010; Taylor & Samuels, 1983).

Als Nebeneffekt der Kenntnis von Textstrukturen resultiert laut Lenhard und Artelt (2009) auch eine höhere Lesegeschwindigkeit, was sich wiederum auf das Leseverständnis auswirkt. Die Lesegeschwindigkeit gilt als Indikator für generelle Lesefähigkeiten und effektives Worterkennen. Durch eine höhere Lesegeschwindigkeit kann mehr kognitive Kapazität für integrative Verständnisprozesse auf Textebene bereitgestellt werden, die bei langsamem, ineffektivem Worterkennen schon auf niedrigeren Ebenen gebunden wäre (Cromley & Azevedo, 2007; Fuchs, Fuchs, Hosp & Jenkins, 2001). So zeigten Bourassa, Levy, Dokin und Casey (1998), dass durch ein Lesegeschwindigkeitstraining signifikant bessere Leistungen in der Lesegeschwindigkeit, der Lesesicherheit (weniger Lesefehler) und im Leseverständnis erzielt werden konnten.

3.3.2.2. Vorwissen

Das Leseverständnis ist von vielen Vorwissensfaktoren des Lesers/der Leserin abhängig. So gelingt die Ausbildung entsprechender Schemata, das Verständnis für Textstrukturen und die Unterscheidung von wichtigen Inhalten und Details nur unter Rückgriff auf entsprechende Wissensbestände. Als entscheidend bei der Ausbildung von Wissensbeständen wird das persönliche Interesse angesehen, das durch unterschiedliche Gründe motiviert sein kann (Klicpera et al., 2010).

Die besondere Funktion des Hintergrundwissens wird im DIME („direct and inferential mediation“) Modell von Cromley und Azevedo (2007) deutlich, in dem ein mittelmäßig direkter Effekt auf das Leseverständnis nachgewiesen werden konnte. Ausführliches Wissen über einen Sachverhalt bedingt ein besseres Verständnis und beeinflusst darüber hinaus die Inferenz- und Strategiebildung, die selbst auch direkt auf das Leseverständnis einwirken.

Der Wortschatz kann in diesem Sinn ebenfalls als Vorwissen aufgefasst werden. Dieses individuell verfügbare Wissen über Wörter stellt einen wesentlichen Einflussfaktor für das Leseverständnis dar (Bundesministerium für Bildung und Forschung BMBF, 2007; Cromley & Azevedo, 2007; Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010; Lenhard & Artelt, 2009, National institute of child health & human development NICHD, 2000; vgl. auch Abschnitt 4.5.). Ein umfangreicher Wortschatz bedingt ein gutes allgemeines Leseverständnis, als Nachweis für diese Annahme gelten unter anderem Trainingsstudien (z.B. Beck, Perfetti & McKeown, 1982). Durch Kenntnis der Wortbedeutung wird ein Text überhaupt erst verstanden. Die Erschließung der Textbasis, der Aufbau propositionaler Strukturen und eines Situationsmodells gelingen leichter, was sich wiederum positiv auf das Verständnis auswirkt (McElvany & Schneider, 2009). Somit kann von einer wechselseitigen Beeinflussung von Leseverständnis und Wortschatz auf allen Ebenen ausgegangen werden.

3.3.2.3. Inferenzen

Die Inferenzbildung bezeichnet den Vorgang des deduktiven Denkens. Beim Lesevorgang ist darunter ein „zwischen oder hinter den Zeilen lesen“ zu verstehen, was als elementare Fertigkeit im Leseverständnisprozess angesehen wird. Um eine kohärente Repräsentation des Textes aufzubauen muss der Leser/die Leserin Schlussfolgerungen (Inferenzen) ziehen, da in Texten Informationen und Zusammenhänge nicht immer explizit mitgeteilt werden (Klicpera et al., 2010; Lenhard & Artelt, 2009). In der Literatur werden verschiedene Formen von Inferenzen unterschieden. So gehen etwa Graesser, Singer und Trabasso (1994) von 13 Inferenztypen aus, die nach dem Grad der Relevanz für das Textverständnis geordnet sind.

Bezüglich des Vorgangs der Inferenzbildung bestehen in der Literatur konträre Ansichten. McCoon und Ratcliff (1992) gehen davon aus, dass Inferenzen, die die lokale Kohärenz sicherstellen und auf unmittelbar verfügbarem Wissen oder expliziten Textaussagen basieren, automatisch gebildet werden. Inferenzen, die die globale Ebene betreffen (z.B. das Handlungsziel), werden hingegen nur beim strategischen, zielbezogenen Lesen gebildet (minimalistische Position). Demgegenüber postuliert die konstruktivistische Position (vgl. Graesser et al., 1994), dass Inferenzen eine wesentliche Komponente des Verstehensprozess sind. Schlussfolgerungen, die das Rezeptionsziel des Lesers/der Leserin, die Kohärenzherstellung auf lokaler und globaler Ebene im referentiellen mentalen Modell, und die Erklärungsversuche für im Text erwähnten Handlungen, Ereignisse und Zustände

betreffen, werden spontan (online) während des Lesens gebildet. In Abhängigkeit vom Leseziel, dem Vorwissen, der Intelligenzleistung (Oakhill & Garnham, 1988) und der Textart stehen dem Leser/der Leserin verschiedene Möglichkeiten offen, einen Text tiefer zu erfassen oder nur oberflächlich zu verarbeiten, wobei allerdings fraglich ist, ob ein nicht-zielbezogenes Lesen (minimalistische Theorie) überhaupt realistisch ist. Die Bildung von Inferenzen scheint darauf hinzuweisen, dass der Leseprozess als ein interaktiver Vorgang zwischen LeserIn und Text zu verstehen ist (Christmann & Groeben, 1999).

3.3.3. Text-Leser-Interaktion: Mentale Modelle

Mentale Modelle verbinden text- und leserorientierte Aspekte beim Textverständnis. Sie werden als „funktionale und strukturelle Analogie zu einem Sachverhalt in der Realität“ definiert (Christmann & Groeben 1999, S. 170). Informationen aus dem gelesenen Text werden stetig mit dem Vorwissen des Lesers abgeglichen und aus diesem Informationsaustausch entsteht ein ständig neu aktualisiertes Situationsmodell.

Eine direkte empirische Überprüfung mentaler Modelle ist nicht möglich, trotzdem betonen Christmann und Groeben (1999), dass es sich bei diesem Theorieansatz um den derzeit besten zur Erklärung der Text-Leser-Interaktion handelt.

4. Entwicklung der Lesefähigkeit

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Die theoretischen Modelle der Leseentwicklung lassen sich bezüglich ihrer Fokussierung auf verschiedene Komponenten des Leseprozesses unterscheiden. Während die Phasenmodelle, die Bezug auf Informationsverarbeitungstheorien (Klicpera et al., 2010) nehmen (Ehri, 1997; Frith, 1985; Günther, 1986) und das Kompetenzentwicklungsmodell (Klicpera et al., 2010) hauptsächlich die Entwicklung des Worterkennens thematisieren, berücksichtigen der „simple view of reading“- Ansatz (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000) und die „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002) zusätzlich das Leseverständnis. Eine Gemeinsamkeit der Theorien besteht in der Annahme, dass die Leseentwicklung bereits vor dem Schuleintritt beginnt.

Diese Modelle stellen die zentralen Aspekte der Leseentwicklung dar. Ziel einer gelungenen Leseentwicklung ist das Textverständnis, das in dieser Studie anhand seiner Teilkomponenten Dekodieren, Hörverstehen und Wortschatz untersucht wird. Es erscheint daher naheliegend zunächst die verwendeten Strategien und Zugänge zum Lesen näher zu betrachten.

Der Vollständigkeit halber sei erwähnt, dass der Schriftspracherwerb stark von Vorläuferfähigkeiten, wie der phonologischen Informationsverarbeitung beeinflusst wird. Der sogenannten phonologischen Bewusstheit, die Fähigkeit, die Struktur der Sprache zu erkennen, wird dabei eine tragende Rolle zugeschrieben (Lenhard, 2005b). Für Elbro (1996) gilt die phonologische Bewusstheit als wichtigster Einzelprädiktor der Leseentwicklung.

Auch die frühe Lesesozialisation innerhalb der Familie, wie unter anderem die Leseförderung und Leseerziehung durch die Eltern oder auch der Umgang mit Literatur innerhalb der Familie beeinflusst die Lesefähigkeit (BMBF, 2007; Steinbrecher, 2007).

4.1. Der Schriftsprachenerwerb nach Frith und Günther

Dieser Abschnitt wurde verfasst von Yvonne Huemer (2011).

Frith (1985) postuliert in ihrem Entwicklungsstufenmodell eine diskrete Abfolge der logographischen, der alphabetischen und der orthographischen Phase (vgl. Graf, 1994; Kirschhock, 2004; Klicpera & Gasteiger-Klicpera, 1998; Klicpera et al., 2010; Lenhard & Artelt, 2009).

Die erste Phase ist durch das unmittelbare Erkennen einzelner bekannter Wörter aufgrund globaler visueller Merkmale gekennzeichnet. In der logographischen Phase werden Wörter in Form von „Logogrammen“ gespeichert, so dass es Kindern gelingt, Reklame-Logos und Artikelbezeichnungen (z.B. McDonalds, Coca Cola) zu identifizieren und in ihrer Bedeutung zu erfassen. Die Anordnung der Buchstaben im Wort und das phonologische Rekodieren sind in diesem Stadium von keinerlei Bedeutung. Durch die zunehmend differenzierte Wahrnehmung der Merkmale (Buchstaben/Buchstabengruppen) wird der Übergang in die nächste Phase eingeleitet.

Die alphabetische Phase ist durch das anwachsende Wissen über Buchstaben und ihre Phonementsprechung sowie über den strukturellen Aufbau und die Funktion der Sprache, charakterisiert. Durch die Fähigkeit zur systematischen Rekonstruktion der Graphem-Phonem-Zuordnung und einer zunehmende Interaktion mit weiteren Teilfertigkeiten, wie etwa dem Dehnlernen, können Fortschritte im Lesen erzielt werden. Das logographemische, direkte Erkennen tritt zunehmend in den Hintergrund. Somit kann die alphabetische Stufe als indirekter Weg gemäß der Zwei-Wege-Theorie aufgefasst werden.

Das orthographische Stadium ist erreicht, wenn durch den Aufbau einer vollständigen inneren Repräsentation der Buchstabenfolge ein direktes Worterkennen möglich wird. Das Lesen eines Wortes führt zur Aktivierung der entsprechenden Information über die Buchstabenfolge im orthographischen Gedächtnis, wobei Wörter vermutlich auch in Morpheme, Silben und häufig vorkommende Buchstabenfolgen gegliedert werden. In Folge der schnellen Wortidentifikation und der kontinuierlichen Automatisierung des phonologischen Rekodierens macht sich eine deutliche Erhöhung der Lesegeschwindigkeit bemerkbar.

Günther (1986) nimmt in seinem Stufenmodell einen fünfstufigen Entwicklungsprozess des Leselernens an. Er geht im Vergleich zu Frith (1985) von einem früheren Beginn des Schriftspracherwerbs aus, der seiner Meinung nach in einer Vorstufe des logographemischen Stadiums, dem präliteralen-symbolischen Stadium, zu sehen ist. Der orthographischen Stufe folgt die integrativ-automatisierte Phase, in der sich die erworbenen Fertigkeiten in einem Prozess der Übung automatisieren und festigen. Erst durch ausreichend Erfahrung im Umgang mit der Sprache wird sinnentnehmendes Lesen möglich (Sassenroth, 2003).

Die Kritik an den Phasenmodellen von Frith (1985) und Günther (1986) betrifft hauptsächlich die Annahme eines längeren logographemischen Stadiums. Insbesondere im deutschen Sprachraum wird dieses, aufgrund seiner regelmäßigen Schriftsprache, angezweifelt. Ein logographemisches Erkennen kann, wenn überhaupt, nur in den ersten Wochen des Erstleseunterrichts beobachtet werden. Unter anderem lassen die Ergebnisse der Wiener Längsschnittstudie (Klicpera & Gasteiger-Klicpera, 1993) darauf schließen, dass durch die im Unterricht vermittelten Kenntnisse über die Phonem-Graphem-Korrespondenz, Wörter eher durch phonologisches Rekodieren erschlossen werden. Schon nach relativ kurzem Leseunterricht sind ErstklässlerInnen in der Lage, unbekannte Buchstabenfolgen sicher zu lesen. Desweiteren ist auch die empirische Überprüfung des Modells als bisher unzureichend zu betrachten (Klicpera & Gasteiger-Klicpera, 1998; Lenhard & Artelt, 2009). Ebenfalls als Nachteil wird die eindeutige Abfolge bestimmter Entwicklungsphasen, in denen jeweils eine bestimmte Strategie zum Tragen kommt, angesehen (Klicpera, Schabmann & Gasteiger Klicpera, 2003).

4.2. Das Modell von Ehri

Dieser Abschnitt wurde verfasst von Maria Klausecker (2011).

Das Modell des Sichtwortlesens von Ehri (1997) ist den Stadienmodellen des Lesens zuzuteilen und bezieht sich dabei explizit auf Informationsverarbeitungstheorien (Klicpera et al., 2003).

In der ersten Phase, in der sogenannten voralphabetischen Phase, werden zum Erkennen einzelner Wörter nur einzelne visuelle Merkmale zum Erkennen der Wörter herangezogen. Anschließend folgen drei alphabetische Phasen (teilweise alphabetische Phase, vollständige

alphabetische Strategie und konsolidierte alphabetische Phase), die den interaktiven Aspekt dieses Modells betonen. Aufeinanderfolgend entwickeln sich in den drei alphabetischen Phasen das Wissen über Buchstaben-Laut-Verbindungen, der sogenannte „Sichtwortschatz“ und der lexikalische Zugang. Dabei wird die Lesegeschwindigkeit sukzessive gesteigert (Ehri, 1997). Ehri (1997) untersuchte ihr Modell in empirischen Studien und konnte eine Bestätigung des Theoriemodells finden.

4.3. Das Kompetenzentwicklungsmodell des Lesens

Dieser Abschnitt wurde verfasst von Yvonne Huemer (2011).

Das Kompetenzentwicklungsmodell des Lesens nach Klicpera et al. (2003) lässt sich unter den Entwicklungsmodellen des Lesens einordnen und berücksichtigt speziell Forschungsergebnisse aus dem deutschen Sprachraum. Die einzelnen unten beschriebenen Entwicklungsphasen laufen dabei nicht strikt nacheinander ab, sondern es werden unterschiedliche Entwicklungsverläufe abhängig von den individuellen Lernvoraussetzungen der Schüler und von der Leseinstruktion berücksichtigt.

Beim geübten Leser/bei der geübten Leserin erfolgt der Worterkennungsprozess entweder direkt über den Zugriff auf das „mentale Lexikon“, in dem die einzelnen Wörter im Gesamten abgespeichert sind, oder über „phonologische Rekodierung“, wobei hier das Wort schrittweise durch die Aufeinanderfolge der einzelnen Buchstaben erlesen wird. Das Erlesen eines Wortes über das mentale Lexikon erfolgt dabei generell schneller. Beide Wege sind beim Leseprozess allerdings unerlässlich. Neue Wörter können nur über die phonologische Rekodierung erlesen werden. Fremdwörter beispielsweise, bei denen sich die Schreibweise stark von der Aussprache unterscheidet, können allerdings nur über den Zugriff auf das mentale Lexikon erlesen werden (Klicpera et al., 2003). Durch starke Interaktion mit der Leseinstruktion entwickeln sich beide beschriebenen Wege des Lesens heraus. Leseinstruktion wird dabei in erster Linie als der Leseunterricht in der Schule verstanden. Allerdings sind auch alle anderen Maßnahmen zur Förderung des Lesens im schulischen Kontext und auch außerhalb der Schule damit gemeint (Klicpera et al., 2003).

Angelehnt an Ehri (1997) wird die erste Phase der Leseentwicklung „präalphabetische Phase“ genannt. Deutschsprechende Kinder haben vor Schuleintritt nur sehr wenige Kenntnisse über die Schrift. Allerdings versuchen die Kinder in dieser Phase Wörter durch einzelne

signifikante Merkmale zu „lesen“. Man kann dies als „logographische“ Phase wie Frith (1985) interpretieren. Unterschiedliche Schwächen und Stärken der Kinder beispielsweise bei der phonologischen Bewusstheit oder bei Gedächtnisprozessen beeinflussen natürlich den Prozess des Lesenlernens und sagen die spätere Leseentwicklung voraus, allerdings nur unvollständig und unter gewissen Bedingungen. Diese Bedingungen sind wie erwähnt abhängig vom jeweiligen Unterricht und der Leseinstruktion. Als kritischer Moment ist der Zeitpunkt der Einschulung anzusehen. Diese Bedingungen sind für individuelle Entwicklungsverläufe maßgeblich beeinflussend (Klicpera et al., 2003).

Die „alphabetische Phase mit geringer Integration“ stellt die erste Phase des tatsächlichen Lesens dar. Die notwendigen Kompetenzen zum Lesen werden schrittweise ausgebildet, wobei sie zu Beginn noch nicht fehlerlos zusammenarbeiten. Zu Beginn des Lesenlernens stehen die „Aneignung des alphabetischen Prinzips und das Erlernen der phonologischen Rekodierung“ (Klicpera et al., 2003, S. 28). Ob bei allen deutschsprechenden Kindern eine „logographische Phase“ nach Frith (1985) auftritt, ist fraglich. Nach dem Modell ist das Auftreten der logographischen Phase abhängig von der Erstleseinstruktion, wobei sie nur bei wenigen Kindern zu beobachten ist und wenn diese auftritt meist auch nur sehr kurz ist (Klicpera et al., 2003).

Die Automatisierung des Lesevorgangs entwickelt sich dann gleichzeitig mit dem phonologischen Rekodieren. Der Aufbau des mentalen Lexikons ist also nicht an eine abgeschlossene Entwicklung des phonologischen Rekodierens gebunden. Spezielle Gedächtnisaspekte und phonologische Rekodierfähigkeiten unterstützen allerdings den Aufbau des mentalen Lexikons (Klicpera et al., 2003).

In der „alphabetischen Phase mit voller Integration“ entwickelt sich die Automatisierung des lexikalischen und auch des nicht-lexikalischen Leseprozesses. Die Lesegeschwindigkeit steigt und es werden weniger Lesefehler begangen. „Partiell lexikalisches Lesen“, also die günstigere Bündelung von Teilprozessen der Informationsverarbeitung, führen zur Zunahme der Lesegeschwindigkeit. Häufig vorkommende Buchstabenfolgen werden als größere schriftsprachliche Einheit abgespeichert und somit schneller erlesen. Auch die Entscheidung für das Erlesen eines Wortes entweder über den lexikalischen oder nicht-lexikalischen Weg wird schneller getroffen und diese beiden Prozesse interagieren stärker. Diese länger

andauernde Phase stellt den Übergang in die letzte „Phase der automatisierten und konsolidierten Integration“ aller beteiligten Verarbeitungsprozesse dar (Klicpera et al., 2003).

Ein Kritikpunkt ist, dass eine einseitige Orientierung des Leseprozesses auf Wortniveau erfolgt. Sinnentnehmendes Lesen, das teilweise bereits bei Leseanfängern ausgebildet ist, kann durch dieses Modell nicht erklärt werden (Heidemann-Menda, 2006).

4.4. Simple view of reading

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Mit dem „simple view of reading“-Ansatz (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000) wurde ein theoretischer Rahmen geschaffen, der den gesamten Lese(lern)vorgang abzubilden versucht. Neben den grundlegenden Kompetenzen der Lesefertigkeit wird dabei auch gleichermaßen das Leseverständnis berücksichtigt (Kirschhock, 2004).

Die Komplexität des Leseprozesses wird von Hoover und Gough (1990) durch die Bezeichnung „simple view“ keineswegs bestritten. Es soll dadurch viel eher verdeutlicht werden, dass lediglich zwei, als gleich wichtig zu betrachtende, (komplexe) Fähigkeiten das Leseverständnis determinieren: das Dekodieren und das Hörverstehen. Für die Annahme, dass das Hörverstehen einen wesentlichen Beitrag zum Leseverständnis leistet, sprechen unter anderem die Befunde von Rost und Hartmann (1992). Bei der Überprüfung ihres hierarchischen Modells des Leseverstehens an einer Stichprobe, bestehend aus 221 Kindern der 4. Grundschulklasse, zeigte sich ein Zusammenhang von $r = .62$ zwischen Hör- und Leseverstehen. Die sprachliche Intelligenz hingegen konnte lediglich 28% der Varianz des Leseverstehens aufklären. Auch Stanovich (z.B. 1989) konnte das Hörverstehen, im Gegensatz zur Intelligenz, als geeigneteren Prädiktor für die Leseleistung identifizieren. Über die Art des Zusammenhangs von Dekodieren und Hörverstehen herrschen unterschiedliche Meinungen. So nehmen Gough und Tunmer (1986) und in späterer Folge auch Hoover und Gough (1990) eine multiplikative Verknüpfung der Komponenten an. Chen und Vellutino (1997) gehen von einem schwächeren und gleichzeitig komplexeren Zusammenhang aus. Allerdings gilt es als empirisch fundiert, dass das Dekodieren und das Hörverstehen

wesentliche Komponenten der Lesefähigkeit beziehungsweise des Leseverständnisses darstellen (vgl. auch Curtis, 1980; Dreyer & Katz, 1992; Sticht & James, 1984).

Das Dekodieren ist in diesem Ansatz mit effizienter Worterkennung gleichzusetzen, was von Hoover und Gough (1990) als “the ability to rapidly derive a representation from printed input that allows access to the appropriate entry in the mental lexicon, and thus, the retrieval of semantic information at the word level” (S. 130) definiert wird. Demzufolge umfasst das Worterkennen die Fähigkeiten Rekodieren (im Sinne eines schnellen Ableitens der Repräsentation des schriftlichen Inputs) und Dekodieren (Zugriff auf das mentale Lexikon und Bedeutungsgenerierung). Ob Wörter im Einzelnen durch buchstabenweises Rekodieren oder auf anderen Weg erkannt werden ist dabei irrelevant. Es wird nur vorausgesetzt, dass ein effizienter (schneller und genauer) Zugriff auf das mentale Lexikon möglich ist. LeseanfängerInnen scheinen sich die Repräsentationen auf einer phonologischen Basis aufzubauen, wobei durch zunehmende Erfahrung phonologische und orthographische Informationen miteinander verbunden werden. Das Rekodieren wird durch effizientere Möglichkeiten zum Worterkennen abgelöst, die nicht näher beschrieben werden. Als Hörverstehen wird von den Autoren die Fähigkeit zur Aufnahme lexikalischer Information (z.B. semantische Information auf Wortebene) bezeichnet, wodurch entsprechende Satz- und Textinterpretationen abgeleitet werden können. Das Leseverstehen betrifft „the same ability, but one that relies on graphic-based information arriving through the eye“ (Hoover & Gough, 1990, S. 131).

Diese Auffassung von modalitätsunabhängigen Verständnisprozessen für das Hör- und Leseverstehen entspricht einer monistischen Sichtweise, die auch für den deutschsprachigen Raum postuliert wird (Marx & Jungmann, 2000; Rost & Buch, 2010; Rost & Hartmann, 1992). Nach dem Worterkennen laufen Verarbeitungs- und Verständnisprozesse für schriftlich oder auditiv dargebotene Informationen grundsätzlich identisch ab. Generell wird davon ausgegangen, dass zu Beginn der Leseentwicklung eine Asymmetrie zugunsten des Verständnisses beim Hören, im Vergleich zum Lesen, besteht. Das Leseverstehen muss folglich im Laufe der Entwicklung dem Hörverstehen angeglichen werden, wobei der Worterkennung (Rekodieren und Dekodieren) entscheidende Bedeutung zukommt. Durch die kontinuierliche Auseinandersetzung mit der Schriftsprache verbessern sich die Fähigkeiten im Worterkennen, wodurch gleichzeitig der Wortschatz, das sprachliche Wissen und die Fertigkeiten im Hörverstehen anwachsen. Verständnisschwierigkeiten entstehen für

LeseanfängerInnen hauptsächlich durch Probleme beim Worterkennen, wohingegen bei geübten LeserInnen die erreichten Fertigkeiten im Hörverstehen die Obergrenze für die Verständnisleistung im Lesen darstellen. Die Beziehung von Hör- und Leseverständnis ist von den erreichten Lesefertigkeiten (Dekodieren) abhängig und die Diskrepanz zwischen den Verständniskomponenten verringert sich mit zunehmend besseren Lesefertigkeiten. Der enge Zusammenhang zwischen Lese- und Hörverstehen (bzw. das Angleichen der beiden Fähigkeiten) bei Erwachsenen/geübten LeserInnen ist auf die zunehmende Automatisierung des Leseprozesses zurück zuführen (Hoover & Gough, 1990; Kirschhock, 2004; Marx & Jungmann, 2000; Rost & Buch, 2010).

In Übereinstimmung mit diesen Hypothesen sei auf eine Reihe von Studien verwiesen. So zeigten Chen und Vellutino (1997), dass die Leistungen im Dekodieren, Lese- und Hörverstehen von der 2. bis zur 6. Klasse kontinuierlich zunahmen. Während die Stärke der Korrelation zwischen Dekodieren und Leseverständnis mit zunehmender Klassenstufe abnimmt, zeigte sich ein umgekehrter Effekt für die Korrelation zwischen Hör- und Leseverstehen. Demzufolge bildet ein gewisses Niveau im Worterkennen die Voraussetzung für Verständnisprozesse. Ist dieses erreicht, wird das Leseverständnis durch das Hörverstehen maßgeblich beeinflusst. Im Gegensatz zu Hoover und Gough (1990), die von einem beginnenden substantiellen Zusammenhang zwischen Lese- und Hörverstehen ab dem Ende der dritten Klasse ausgehen, konnte Hagtvat (2003) bereits für 9-jährige Kinder aus Norwegen am Ende der zweiten Schulstufe einen stärkeren Einfluss des Hörverstehens, im Vergleich zum Dekodieren, auf das Leseverständnis nachweisen. Eine mögliche Schlussfolgerung wäre, dass in regelmäßigen Schriftsprachen (z.B. Norwegisch, Deutsch) das Worterkennen durch die einfache Graphem-Phonem-Korrespondenz früher im Leselernprozess automatisiert wird. de Jong & van der Leij (2002) interpretierten ihre Forschungsergebnisse ebenfalls im Sinne einer Bestätigung des „simple view of reading“- Ansatzes. Sie zeigten, dass die Geschwindigkeit des Worterkennungsprozesses und das generelle sprachliche Verständnis die Entwicklung des Leseverständnisses von der ersten bis zur dritten Schulstufe beeinflussen und dass dieser Einfluss stabil über diese drei Jahre bleibt.

Vergleichbare Ergebnisse konnten von Rost und Hartmann (1992) für den deutschsprachigen Raum beobachtet werden, wobei insbesondere die Replikationsbefunde von Marx und Jungmann (2000) erwähnenswert sind. In ihrer Untersuchungsstichprobe von 360 Kindern stiegen die Fähigkeiten im Hör- und Leseverstehen über alle sieben Messzeitpunkte (vom

Ende der ersten Klasse bis Ende der vierten Klasse) hinweg an, wobei die Leistungen im Hörverstehen stets besser waren und die obere Leistungsgrenze im Leseverstehen bildeten. Die theoretisch postulierte Annäherung zwischen Hör- und Leseverständnis zeigte sich für den gesamten Untersuchungszeitraum und am Ende der vierten Klasse war eine Diskrepanz zwischen den beiden Verstehenskomponenten kaum noch beobachtbar. Dies ist durch die enormen Zuwachsraten im Leseverstehen bei vergleichsweise geringerem Anstieg des Hörverstehens bedingt, wobei bemerkenswert ist, dass sich die Leistungen im Leseverständnis bis zum Ende der dritten Klasse mehr als verdoppelt haben. Diese Entwicklung wird auf die zunehmende Automatisierung des Worterkennens zurückgeführt und dieser Effekte kann bis in die sechste Klasse nachgewiesen werden. Die Relevanz des Hörverstehens im Leseverständnisprozess lässt sich bereits ab Mitte der zweiten Klasse nachweisen ($r = .60$ bis $r = .65$).

Basierend auf den Annahmen des „simple view of reading“-Ansatz entwickelten Marx und Jungmann (2000) ein Modell des Leselernens. Im Unterschied zu den amerikanischen Modellen, wird neben den vorschulischen Bedingungsvariablen auch der indirekte Einfluss des Hörverstehens auf die Entwicklung grundlegender Lesefertigkeiten (Worterkennen) berücksichtigt. Das Rekodieren, als Aussprache des Geschriebenen ohne Bedeutungserfassung definiert, gilt neben dem Eintrag des Wortes im mentalen Lexikon als Voraussetzung für das Dekodieren (Erfassen der Bedeutung).

Obwohl durch den „simple view of reading“-Ansatz die wichtige Rolle des Dekodierens und des Hörverstehens für die Entwicklung des Lesens und deren Einfluss auf das Leseverständnis überzeugend belegt sind, wurden Erweiterungen und Modifikationen angeregt. Zum einen wird die zusätzliche Berücksichtigung der Dekodiergeschwindigkeit (fluency) angeregt, da für diese ein signifikanter Beitrag zur Varianzaufklärung im Leseverständnis nachgewiesen werden konnte (Joshi & Aaron, 2000). Zum anderen belegen die Ergebnisse von Braze, Tabor, Shankweiler und Mencl (2007), dass der Wortschatz als eigenständige Komponente in das Modell eingefügt werden sollte. Der Wortschatz scheint nicht nur für die Ausbildung eines umfassenderen Hörverstehens verantwortlich zu sein, sondern auch einen wesentlichen Beitrag zum Leseverständnis zu leisten (zusätzliche Varianzaufklärung von 6%). Indem Tilstra, McMaster, van der Broek, Kendeou und Rapp (2009) das ursprüngliche Modell des „simple view of reading“-Ansatzes durch die Dekodiergeschwindigkeit und den Wortschatz/

verbale Fertigkeiten ergänzten, konnten beachtliche 74% der Varianz des Leseverständnisses bei amerikanischen SchülerInnen der 4. Klasse aufgeklärt werden.

4.5. Lexical quality hypothesis

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Ein Ansatz, der explizit die zentrale Bedeutung von Wortkenntnissen im Leseverständnisprozess berücksichtigt, wurde von Perfetti und Hart (2001, 2002) konzipiert. Der „lexical quality hypothesis“ (LQH) zufolge beruht ein Großteil des Leseverständnisses auf dem detaillierten Wissen über Wörter, das durch individuelle Erfahrungen mit gesprochener und geschriebener Sprache erworben wird.

Die Grundlage der LQH bildet die bereits 1985 von Perfetti entwickelte „verbal efficiency theory“. Deren Hauptaussage besteht darin, dass das Leseverständnis von erfolgreichem Worterkennen (i.S. eines schnellen Abrufs der Phonologie und der Wortbedeutung) abhängig ist und Unterschiede in der Verständnisleistung durch unterschiedlich ausgeprägte Fähigkeiten im Wortlesen bedingt sind. Verarbeitungsprozesse auf Wortebene laufen zunehmend schneller und automatischer ab, sodass weniger kognitive Verarbeitungsressourcen benötigt werden, die dann für komplexere, hierarchiehöhere Vorgänge zur Verfügung stehen. Diese enge Beziehung zwischen Dekodierfähigkeit und Leseverständnis ist in zahlreichen Untersuchungen, auch für geübte LeserInnen, überzeugend belegt worden (z.B. Bell & Perfetti, 1994; Shankweiler et al., 1999; Tilstra et al., 2009). Allerdings lässt sich durch diese Annahme nicht hinreichend erklären, welche Faktoren im Wortverarbeitungsprozess die Variation im Verständnis bedingt, da Effizienz nicht mit Geschwindigkeit gleichzusetzen ist. Auch Perfetti (2007) zufolge ist die Fähigkeit, Wörter zu identifizieren und deren Bedeutung in einem bestimmten Kontext zu erfassen ausschlaggebend, wobei er jedoch zusätzlich annimmt, dass „this source of this ability is the knowledge a reader has about words“ was als „specific lexical representations“ (S. 359) bezeichnet wird.

Das Wortwissen ist in einem (mental) Lexikon gespeichert und besteht aus orthographischen, phonologischen sowie semantischen Spezifikationen eines Wortes. Der semantische Bestandteil der lexikalischen Repräsentation umfasst neben der Wortbedeutung

auch die grammatikalischen Informationen. Die drei Konstituenten sind in Form eines Netzwerkes miteinander verbunden, sodass die Wortidentifikation durch den Zugriff auf die lexikalische Repräsentation möglich wird (Perfetti, 2007; Perfetti & Hart, 2002). Die exakte Kenntnis von Wortformen (Schreibweise, Aussprache, Grammatik) und Wortbedeutungen wird als lexikalische Qualität („lexical quality“; LQ) bezeichnet, die bei entsprechend hoher Ausprägung schnelle Verarbeitungsprozesse ermöglicht (Perfetti & Hart, 2001). Die Struktur einer lexikalischen Repräsentation ist in dem Ausmaß als gut ausgeprägt („high quality“) zu sehen, in dem „it [lexical representation] has a fully specified orthographic representation (spelling) and redundant phonological representations (one from spoken language and one recoverable from orthographic-to-phonological mappings)“ (Perfetti & Hart, 2002, S. 190).

Als grundlegende Annahme der LQH wird postuliert, dass die Lesefähigkeit auf einer reliablen und kohärenten Repräsentation der Bestandteile beruht. Eine lexikalische Repräsentation ist in dem Sinn als zuverlässig und stabil aufzufassen, wenn trotz korrespondierender Schreibweise oder Aussprache das vorgegebene Wort korrekt identifiziert werden kann. Die Kohärenz wird durch die Verbindung der Konstituenten gesichert, wodurch diese simultan für eine effiziente Wortidentifikation verfügbar sind und vermehrte kognitive Ressourcen für die Bedeutungserfassung auf Textebene frei werden. Desweiteren gehen Perfetti und Hart (2002) davon aus, dass sich die lexikalischen Repräsentationen durch Erfahrungen im Umgang mit der Sprache ausbilden. Zu Beginn der Leseentwicklung unterstützen phonologisches Wissen und die Fähigkeit zum Rekodieren die Entstehung wortspezifischer orthographischer Spezifikationen. Durch die kontinuierliche Auseinandersetzung mit der Sprache beim Lesen, Schreiben, Hören und Sprechen entstehen detaillierte und neue Repräsentationen. Das Ausmaß an sprachlicher Erfahrung beeinflusst die lexikalische Qualität für ein bestimmtes Wort, sowie die Gesamtanzahl von Wörtern mit hoher/geringer Qualität. Lexikalische Fertigkeiten bewirken demnach ein besseres (Lese-) Verständnis, was gemäß Stanovich (1986; Matthäus-Effektes in der Leseentwicklung) zu mehr Lesepraxis führt, die wiederum die lexikalischen Fertigkeiten verbessert. Dieser sich selbst verstärkende Effekt verdeutlicht den Autoren zufolge die enge Beziehung zwischen den lexikalischen Komponenten und dem Verständnis.

Die Effekte lexikalischer Qualitäten auf das Verständnis wurden in einer Reihe von Experimenten untersucht. Durch die Vorgabe von Homophonen (Wörter mit gleicher Phonologie und unterschiedlicher Bedeutung, z.B. „gait“ und „gate“, vgl. Perfetti & Hart,

2002, S. 193) wurde die angenommene Reliabilität überprüft. Es zeigte sich, dass sowohl gute wie auch schlechte LeserInnen durch die Homophone in ihrem Wortverarbeitungsprozess gestört werden. Schlechte LeserInnen scheinen insgesamt langsamer in ihrer Verarbeitung zu sein, was auf die schwachen Verbindungen im Netzwerk zurückgeführt wird. Ein weiterer Aspekt betrifft die Tatsache, dass die Häufigkeit, mit der ein Wort vorkommt, entscheidend für dessen Verarbeitung ist. Gute LeserInnen, die der LQH entsprechend über mehr Erfahrung mit der Schriftsprache verfügen und auf stabilere lexikalische Repräsentationen zurückgreifen können, zeigten kaum Irritationen beim Lesen des häufiger vorkommenden Homophon-Wortes. Schlechten LeserInnen hingegen sind auch häufigere Wörter eher unbekannt und diese führen folglich zu Irritationen, da durch die weniger stabilen Repräsentationen beide Wörter aktiviert werden. Desweiteren wurde die angenommene Notwendigkeit kohärenter lexikalischer Repräsentationen zur effizienten Wortidentifikation faktorenanalytisch bestätigt. Zusammenfassend lässt sich folgern, dass lexikalische Repräsentationen von hoher Qualität die Genauigkeit und Geschwindigkeit der Worterkennung, die allgemeinen Lesekompetenzen, die Fähigkeit zum Erlernen neuer Wörter und ihrer Bedeutungen, sowie letztendlich das Leseverständnis positiv beeinflussen (Perfetti, 2007; Perfetti & Hart, 2002; Shankweiler, Lundquist, Dreyer & Dickinson, 1996).

In Studien zum Thema Leseverständnis wird häufig der Wortschatz als Indikator für lexikalische (bzw. verbale) Fertigkeiten herangezogen. Obwohl dieser streng genommen nur den semantischen Konstituenten der lexikalischen Repräsentation erfasst, kann gleichzeitig argumentiert werden, dass sich ein detaillierter Wortschatz durch die Verbindung der phonologischen, orthographischen und semantischen Bestandteile aufbaut, was in erster Linie durchs Lesen geschieht. Somit kann davon ausgegangen werden, dass der Wortschatz einen geeigneten Index für umfassende lexikalische Qualitäten darstellt (Protopapas, Sideridis, Mouzaki & Simos, 2007). Die Ergebnisse von Braze et al. (2007) und Tistra et al. (2009) betonen den eigenständigen, wichtigen Anteil der lexikalischen Fertigkeiten für das Leseverständnis. Auch Protopapas et al. (2007) konnten einen direkten Effekt des Wortschatzes auf das Leseverständnis für griechische VolksschülerInnen bestätigen. Die Effekte der Dekodierfähigkeit (Genauigkeit und Geschwindigkeit) wurden indirekt über den Wortschatz auf das Leseverständnis vermittelt. Yovanoff, Duesbery, Alonzo und Tindal (2005) identifizierte ebenfalls den Wortschatz und die mündliche Leseflüssigkeit als relevante Determinanten des Leseverständnisses (vgl. auch Abschnitt 3.3.2.2.)

Das National Institute of Child Health and Human Development (NICHD, 2000) definiert den Wortschatz als einen der wichtigsten Prädiktoren für den Erfolg im Leseprozess. Auch Stahl und Fairbanks (1986) konnten den Zusammenhang zwischen Lesekompetenz und dem Wortschatz aufzeigen. Trainings im Bereich des Wortschatzes wirken sich zudem positiv auf das Textverständnis aus. Der Wortschatz zeigt allerdings positive Auswirkungen sowohl auf Wort-, Satz- und Textebene. McElvany und Becker (2007, zitiert nach Lenhard & Artelt, 2009) konnten in einer Längsschnittstudie an Berliner Schulen mit SchülerInnen der dritten bis zur sechsten Schulstufe aufzeigen, dass zwischen der Wortschatz- und der Textverständnissentwicklung eine wechselseitige Beeinflussung herrscht. Auf der einen Seite erlaubt ein umfassender Wortschatz einen schnelleren und sicheren Zugriff auf das mentale Lexikon (NICHD, 2000). Auf der anderen Seite haben Kinder mit geringem Leseverständnis größere Probleme die Bedeutung von unbekannten Wörtern zu erschließen. Die Wortschatzentwicklung wird dadurch verzögert (Cain, Oakhill & Lemmon, 2004).

Empirischer Teil

5. Zielsetzung und Fragestellungen

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

5.1. Zielsetzung

Das Ziel der drei Parallelstudien dieses Forschungsprojektes ist eine Überprüfung der kausalen Zusammenhänge zwischen einzelnen Prädiktorvariablen und dem Leseverständnis bei Kindern der zweiten (Klausecker, 2011), dritten (Huemer, 2011) und vierten Grundschulstufe. In der vorliegenden Diplomarbeit wird angestrebt, den direkten Einfluss des Worterkennens, des Wortschatzes und des Hörverstehens auf das Leseverständnis bei ViertklässlerInnen zu erfassen, sowie die Beziehung der Prädiktoren untereinander zu klären.

Obwohl in zahlreichen Studien der Einfluss des Dekodierens und des Hörverstehens auf das Leseverständnis überzeugend belegt werden konnte (vgl. Abschnitt 4.4.) und auch dem Wortschatz eine wesentliche Rolle im Leseverständnisprozess zugesprochen wird (vgl. Abschnitt 3.3.2.2 und Abschnitt 4.5.), fehlt bisher eine empirische Fundierung für ein umfassendes Modell des Leseverständnisses, insbesondere für den deutschsprachigen Raum. Durch die vorliegende Untersuchung soll festgestellt werden, ob und in welchem Ausmaß jede der Variablen *Worterkennen*, *Hörverstehen* und *Wortschatz* eigenständig zur Varianzaufklärung im *Leseverständnis* beitragen kann. Es wird versucht den „simple view of reading“-Ansatz (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000) und die „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002) in ein gemeinsames Modell zu integrieren.

In Anlehnung an diesen theoretischen Rahmen werden den SchülerInnen vier standardisierte Testverfahren vorgelegt, um die Variablen zu operationalisieren. Es werden die Leistungen in einem Textverständnistest (ELFE 1-6; Lenhard & Schneider, 2006), einem Hörverständnistest (Knuspel-L; Marx, 1998), einem Wortschatztest (adaptierte Version des PPVT-III, urspr. Dunn & Dunn, 1997) und einem Wortlesetest (Schabmann, Schmidt, Klicpera, Gasteiger-Klicpera & Klingebiel, 2009) erfasst. Die Zusammenhänge zwischen den einzelnen Variablen und ihr Einfluss auf das Leseverständnis werden anschließend mittels

Strukturgleichungsmodellen von den Autorinnen für die jeweilige Klassenstufe gesondert analysiert, wobei diese im Diskussionsteil interpretiert und kritisch hinterfragt werden sollen. Der Diskussionsteil umfasst ebenfalls eine Ausführung zum Thema Entwicklungsaspekte des Leseverständnisses bei Grundschulkindern. Auch mögliche praktische Implikationen der Befunde sollen angedeutet werden.

5.2. Fragestellungen

Innerhalb dieses Gesamtforschungsprojektes wird in dieser Diplomarbeit auf Basis des vorangegangenen Theorieteils die Fragestellung, ob das *Leseverständnis* bei Kindern der vierten Volksschulstufe durch die Komponenten des „simple view of reading“- Ansatzes (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000), *Worterkennen* und *Hörverstehen*, sowie durch den *Wortschatz* (in Anlehnung an die „lexical quality hypothesis“; Perfetti & Hart, 2001, 2002, auch Protopapas et al., 2007) adäquat abgebildet werden kann.

Abbildung 1 verdeutlicht die angenommene kausale Beeinflussung des Leseverständnisses durch die Variablen *Worterkennen*, *Wortschatz* und *Hörverstehen*. Das Worterkennen ist durch die beiden Komponenten *Lesegenauigkeit (accuracy)* und *Leseflüssigkeit (fluency)* bedingt, die anhand von Wörtern und Pseudowörtern erhoben werden.

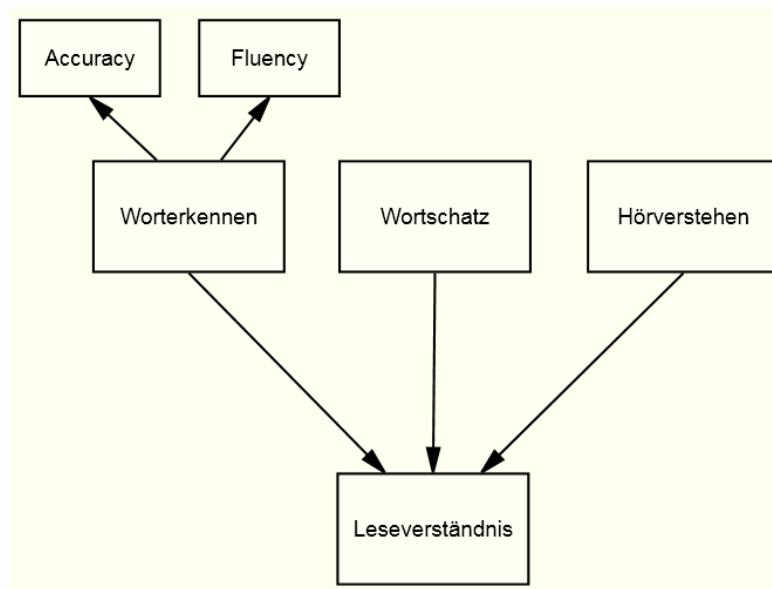


Abbildung 1: Schematische Darstellung der Modelle des Leseverständnisses

Ausgehend von diesen Grundannahmen sollen nun drei verschiedene Modelle bezüglich ihrer Prädiktion des Leseverständnisses überprüft werden. Für Modell 1 wird angenommen, dass das *Worterkennen*, der *Wortschatz* und das *Hörverstehen* jeweils einen direkten Effekt auf das *Leseverständnis* besitzen. Bedingt durch den Schriftsprachenerwerb und der damit einhergehenden zunehmenden Erfahrung mit der Schriftsprache entwickeln sich die Fähigkeiten im Worterkennen bei SchülerInnen rasant. Da Wörter zunehmend sicherer und effizienter erlesen werden, stehen vermehrt kognitive Ressourcen für höhere Verständnisprozesse (z.B. Erfassen von Aussagen auf Satz- und Textniveau) zur Verfügung. Im Zuge dessen findet eine Erweiterung des Wortschatzes statt und somit eine Verbesserung der Fertigkeiten im Hörverstehen (Marx & Jungmann, 2000), sodass ein positiver Zusammenhang zwischen den Variablen angenommen wird. Die „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002) geht davon aus, dass sich lexikalische Repräsentationen durch die Erfahrung mit der Schriftsprache und der mündlichen Sprache ausbilden. Sind die Wörter kohärent und reliabel repräsentiert, dann wäre zu vermuten, dass die lexikalischen Repräsentationen mit der Worterkennung und dem Hörverstehen korrelieren und einen direkten Einfluss auf das Leseverständnis aufweisen sollten (Perfetti, 2007, Perfetti & Hart, 2002). Dementsprechend wird für Modell 1, ausgehend von Abbildung 1, die zusätzliche Interaktion der drei Prädiktorvariablen postuliert.

Perfetti und Hart (2002) zufolge erleichtert ein großer Wortschatz das Worterkennen durch seine gefestigten lexikalischen Repräsentationen. Anhand von Modell 2 soll, im Unterschied zu Modell 1, die Annahme überprüft werden, dass der Wortschatz einen direkten Effekt auf das Worterkennen ausübt. Die für Modell 1 angenommenen direkten Pfade und Korrelationen zwischen den übrigen Variablen bleiben in unveränderter Form bestehen (siehe Abbildung 1).

Die Hypothese, dass die Fähigkeiten im Worterkennen den Wortschatz kausal beeinflussen (vgl. Perfetti & Hart, 2002; Protopapas et al., 2007), wird in Modell 3 überprüft. Weiterhin werden direkte Effekte von Worterkennen, Wortschatz und Hörverstehen auf das Leseverständnis vermutet, sowie positive Beziehungen zwischen Hörverstehen, Wortschatz und Worterkennen.

Durch die im Rahmen dieser Untersuchung aufgestellten Modelle lassen sich lediglich Aussagen über das Leseverständnis in Abhängigkeit von den drei Prädiktorvariablen treffen.

Die angenommen Beziehungen der Variablen untereinander sind, sofern nicht in den verschiedenen Modellen direkte Pfade zwischen dem Worterkennen und dem Wortschatz angenommen werden, ausschließlich korrelativer Art und auch der Einfluss des Leseverständnisses auf das Worterkennen, den Wortschatz und das Hörverstehen wird nicht geklärt.

6. Methodik

6.1. Untersuchungsdurchführung

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Die Untersuchung fand an vier verschiedenen Schulen (Bezirk Vöcklabruck) in ländlichen Gebieten Oberösterreichs von Dezember 2010 bis Februar 2011 statt. Die Studie wurde den DirektorInnen der einzelnen Schulen kurz vorgestellt und ihr Einverständnis zur Teilnahme wurde eingeholt. Vor Beginn der eigentlichen Erhebung wurde um Bewilligung der Studie beim Landesschulrat für Oberösterreich ersucht. Das Ansuchen an den Landesschulrat und die Bewilligung der Studie sind im Anhang zu finden (siehe Anhang A, Anhang B).

Nachdem die Durchführung der Studie bewilligt worden war, wurden Elternbriefe (siehe Anhang C) ausgegeben. Diese Briefe umfassten eine kurze Beschreibung der angewendeten Testverfahren, Aufklärungen zum Datenschutz sowie Angaben zur Kontaktaufnahme bei auftretenden Rückfragen. Die Einverständniserklärung zur Teilnahme an der Studie war zum Schluss auszufüllen.

Als die Teilnahmeerklärungen vollständig retourniert waren, wurde mit der Datenerhebung begonnen. Bei der Datenerhebung wurde eine hinreichend große Stichprobe pro Klassenstufe ($N > 100$; Buch, 2007) zur Analyse mittels Strukturgleichungsmodellen angestrebt. Die Daten wurden an einem Testzeitpunkt im schulischen Kontext der Kinder gewonnen und die Testbedingungen konnten für alle TeilnehmerInnen gleichermaßen realisiert werden.

Zu Beginn der Untersuchung wurden den Kindern zunächst Ziel und Vorgehensweise erläutert und die demographischen Daten erfasst. Die Erhebung beanspruchte pro Klasse zwei Schulstunden, wobei zuerst eine Gruppentestung und dann eine Einzeltestung durchgeführt wurden. Die nach der standardisierten Instruktion begonnene Gruppentestung erstreckte sich über eine Unterrichtsstunde. Es wurde zunächst der Untertest „Hörverstehen“ aus der Testbatterie „Knuspels Leseaufgaben“ (Knuspel-L, Marx, 1998) vorgegeben. Anschließend wurde das Leseverständnis durch die Vorgabe des Subtests „Textverständnis“ aus dem „Leseverständnistest für Erst- bis Sechstklässler“ (ELFE 1-6, Lenhard & Schneider, 2006)

erfasst. Den Abschluss der Gruppenerhebung bildete die von den Autorinnen vorgenommene adaptierte, 15-minütige Version des „Peabody Picture Vocabulary Test“ (PPVT-III, Dunn & Dunn, 1997). Nach einer Pause wurde der Wortlesetest (Schabmann et al., 2009) jedem Kind einzeln zur Bearbeitung vorgegeben. Jeder Testvorgabe gingen eine gemeinsame Bearbeitung von Testitems mit den SchülerInnen und die Klärung anstehender Fragen voraus.

Die Gruppenerhebung wurde im jeweiligen Klassenzimmer durchgeführt. Die Einzeltestung wurde entweder in einem gesonderten Raum oder innerhalb der Schulstunden am Gang durchgeführt. So konnte für alle Kinder eine ruhige und gleiche Testatmosphäre geschaffen werden. Die Rückmeldung der Ergebnisse erfolgte aus Datenschutzgründen über eine Gesamtrückmeldung über alle teilnehmenden Schulen hinweg.

6.2. Erhebungsinstrumente

Bei der Datenerhebung wurden vier verschiedene Tests vorgegeben. Dabei handelt es sich um den Untertest Hörverstehen aus dem Knuspel-L (Marx, 1998), um den Textverständnistest aus dem ELFE 1-6 (Lenhard & Schneider, 2006), um den PPVT-III (Dunn & Dunn, 1997) sowie um den Wortlesetest von Schabmann und Kollegen (2009). Die Vorgabe der Tests erfolgt in der eben genannten Reihenfolge (siehe Anhang D).

6.2.1. Knuspels Leseaufgaben (Knuspel-L)

Dieser Abschnitt wurde von Yvonne Huemer (2011) verfasst.

Der von Marx (1998) konzipierte Lesetest stellt ein ökonomisches und kindgerechtes Verfahren dar, das die Abbildung der Leseleistung von Kindern im Grundschulalter ermöglicht. Mit dem Knuspel-L können die für das verstehende Lesen notwendigen Vorläuferfertigkeiten und die grundlegenden Lesefertigkeiten ganzheitlich erfasst werden.

Das Leseentwicklungsmodell, das der Testkonstruktion zugrunde liegt, ist dem Autor zufolge als Erweiterung und Integration des „simple view of reading“- Ansatzes (Gough & Tunmer, 1986; Hoover & Gough, 1990), der Stufenmodelle (z.B. Ehri, 1997; Frith, 1985; Günther, 1986) und der konnektionistischen Leselernmodelle (z.B. Seidenberg & McClelland, 1989) zu sehen. Neben den spezifischen Lesefertigkeiten werden das Hörverstehen und die

unterschiedlichen Graphem-Phonem-Korrespondenzen mit einbezogen. Die phonologische Route wird als Ausgangspunkt für das Leseverstehen betrachtet.

Das Lesen wird als Fähigkeit zum Rekodieren, Dekodieren und Verstehen schriftsprachlicher Äußerungen verstanden, wobei als basale Lesefertigkeit das Identifizieren von Wörtern ohne Erfassen des Wortsinns verstanden wird (Rekodieren). Darauf aufbauend und unter Verwendung orthographischen Wissens, werden Wörter in ihrem Bedeutungsgehalt erfasst (Dekodieren). Auf dieser Wortebene wird bereits ein Beitrag des semantischen Hörverstehens für die Verstehensleistung angenommen. Schlussendlich werden Sinnzusammenhänge auf Satz- und Textebene hergestellt. Dieser Vorgang wird als vom semantischen (Welt-) und prozeduralen Wissen beeinflusst, und vom Hörverstehen abhängig, vermutet.

In Knuspels Leseaufgaben werden die drei Aspekte der Lesefähigkeit und das Hörverstehen durch entsprechende Subtests operationalisiert. Durch die Annahme einer modularitätsunabhängigen Informationsverarbeitung unterscheiden sich die Untertests zum Hör- und Leseverstehen nur in ihrer jeweiligen Darbietungsform (mündlich vs. schriftlich). Der Versuchsperson wird eine Instruktion erteilt und die gezeigte Reaktion bildet dementsprechend die Bewertungsgrundlage. Es wird davon ausgegangen, dass eine korrekte Ausführung nur dann möglich ist, wenn die Instruktion verstanden wurde.

Der Knuspel-L liegt in zwei Paralleltestformen (Form A und Form B) vor und kann als Gruppen- oder Einzeltestung durchgeführt werden, wobei die Zeitbeschränkung für die Aufgabenbearbeitung in Abhängigkeit der Schulstufe variiert. Die vier Untertest sind:

- Subtest 1-Hörverstehen: Verstehen mündlich gestellter Fragen und Aufforderungen
- Subtest 2-Rekodieren: Erkennen von lautgleichen Wörtern
- Subtest 3-Dekodieren: Erkennen von Wortbedeutungen
- Subtest 4-Leseverstehen: Verstehen schriftlich gestellter Fragen und Aufforderungen.

Die Gütekriterien des Knuspel-L sind als zufriedenstellend zu bewerten (siehe Marx, 1998). Die Reliabilität der einzelnen Subtests erreicht für ein zwölf monatiges Testintervall (Gesamtstichprobe N= 4746) Werte von $r = .620$ bis $r = .791$ (Retestreliabilität). Die

Paralleltestreliabilität ist ebenfalls im mittleren bis hohen Bereich anzusiedeln ($r = .605$ bis $r = .808$). Desweiteren wurde der Knuspel-L anhand von LehrerInnenurteilen, anderen Lesetests und den Schulnoten der Kinder validiert (Werte für die 4. Schulstufe von $r = .292$ bis $r = .564$). Die Objektivität wird durch detaillierte Angaben im Testhandbuch zur Instruktion, Durchführung, Auswertung und Interpretation sicher gestellt, wobei Lenhard (2005a) kritisch anmerkt, dass die Auswertungsobjektivität durch die individuellen Eigenschaften des Kindes bei der Aufgabenbearbeitung negativ beeinflusst sein kann.

Im Zusammenhang mit der vorliegenden Studie wird Subtest 1-Hörverstehen (Anhang D) als Instrument zur Erfassung dieser Fähigkeit ausgewählt. Aufgrund der zufriedenstellenden Testgütekriterien, der mangelnden Verfügbarkeit ähnlich ökonomischer Verfahren im deutschsprachigen Raum und der adäquaten Abbildung des Hörverstehens ohne Konfundierung durch andere Variablen (z.B. Leseverstehen), scheint dieses Vorgehen berechtigt.

Die insgesamt 14 Aufgaben werden den Kindern mündlich vorgetragen, wobei durch die auditive Informationsaufnahme entsprechendes Wissen aktiviert und in korrekte Handlungen umgesetzt werden muss. Die Items enthalten Wissens- und Ausführungsaspekte, die sich entweder auf die eigene Person und Testsituation, oder auf die „Knuspel“ (Fabelwesen des Testverfahrens) beziehen. Die Bearbeitungsdauer pro Item unterscheidet sich je nach intendierter Schwierigkeit. Den Kindern werden die Aufgabenstellungen nur einmal laut vorgelesen. Sie müssen genau zuhören und anschließend ihr Wissen auf Basis der vorgelesenen Instruktion aktivieren und dieses dann in korrekter Weise ausführen (Marx, 1998).

Beispiel (Marx, 1998): „In welchem Monat hast du Geburtstag? Schreibe nur die ersten drei Buchstaben des Monats in Druckbuchstaben in den mittleren Teil des Kästchens“ (Zeit: 20 Sekunden)

Für die statistische Auswertung wird ein Gesamtscore berechnet. Dieser ergibt sich gemäß Marx (1998) aus der Addition der einzelnen Rohpunkte, die für richtig bearbeitet Teilaspekte (mit 1 bewertet) pro Item vergeben werden. Die Items betreffen in unterschiedlicher Weise

Wissen- und/oder Ausführungsdimensionen, sodass insgesamt ein Score von 34 Rohwertpunkten bei vollständig korrekter Aufgabenbearbeitung erzielt werden kann.

6.2.2. Ein Leseverständnistest für Erst- bis Sechstklässler (ELFE 1-6)

Dieser Abschnitt wurde von Maria Klausecker (2011) verfasst.

Der ELFE 1-6 (Lenhard & Schneider, 2006) ist ein normiertes, zuverlässiges, valides und ökonomisches Testverfahren für SchülerInnen der ersten bis sechsten Schulstufe und erlaubt eine umfassende Bewertung des Leseverständnisses auf unterschiedlichen Ebenen. Die Autoren sehen das Leseverständnis als eine Fähigkeit an, die auf verschiedenen Teilfertigkeiten und Bedingungsfaktoren beruht. Die Verständniskomponenten und –prozesse können auf Wort-, Satz und Textebene unterschieden werden (vgl. Christmann & Groeben, 1999, bzw. Abschnitt 3.), wobei auch eine wechselseitige Beeinflussung der relevanten Teilkomponenten angenommen wird.

Von diesem theoretischen Hintergrund ausgehend wird die Erhebung des Leseverständnisses durch drei Subtests (in zwei Paralleltestformen) realisiert:

- Der Wortverständnistest erfasst die basalen Lesefertigkeiten auf Wortniveau (Dekodieren und Synthese).
- Der Satzverständnistest überprüft die syntaktischen Fähigkeiten und das sinnentnehmende Lesen auf dieser Stufe.
- Der Textverständnistest zielt auf die Erhebung der Fähigkeiten zum satzübergreifenden Lesen, schlussfolgernden Denken und Auffinden von Informationen ab.

Der ELFE 1-6 kann als Gruppen- oder Einzeltestung durchgeführt werden und liegt zusätzlich in computerisierter Form vor (zusätzliche Messung der Lesegeschwindigkeit).

Bezüglich der Durchführung, Auswertung und Interpretation kann der ELFE 1-6 als in hohem Maße objektiv bezeichnet werden. Die Items der einzelnen Untertests erfassen jeweils in befriedigendem Ausmaß die vermutete Dimension, die Werte der inneren Konsistenz liegen

für die Subtests zwischen $r = .92$ und $r = .97$. Die Ergebnisse des ELFE 1-6 sind als zuverlässig zu betrachten, die Retestreliabilität des Gesamttests (für SchülerInnen der 4. Klasse) liegt bei $r = .947$. Auch kann davon ausgegangen werden, dass die beiden Testformen äquivalent sind (im Mittel $r = .873$). Die Korrelationen der Ergebnisse des ELFE 1-6 mit den Lehrerurteilen ($r = .705$) und verschiedenen anderen Testverfahren (z.B. „Würzburger Leise LeseProbe“; WLLP: $r = .710$) lassen auf eine ausreichende Testvalidität schließen (Lenhard & Schneider, 2006).

Aus ökonomischen Gründen und aufgrund vorliegender Testgütekriterien wurde für diese Erhebung der Textverständnistest (Form A; Anhang D) als Index für das Leseverständnis ausgewählt. Den SchülerInnen werden kurze, narrative Texte mit zugehörigen Fragen präsentiert. Aus vier Antwortmöglichkeiten ist die richtige Alternative zu wählen. Die insgesamt 20 Items verlangen unterschiedliche Fähigkeiten. So erfordern fünf Aufgaben eine isolierte Informationsentnahme, bei acht Items müssen anaphorische Bezüge hergestellt werden und bei sieben Items ist zur Lösung die Bildung von Inferenzen notwendig. In der auf sieben Minuten befristeten Bearbeitungszeit sollen selbständig so viele Aufgaben wie möglich gelöst werden, wobei die Kinder still an ihren Aufgaben arbeiten (Lenhard & Schneider, 2006).

Beispiel (Lenhard & Schneider, 2006): Paula ist mit ihren Eltern in den Ferien ans Meer gefahren. Am Strand spielt sie im Sand und sammelt schöne, farbige Muscheln. Die findet sie schön.

Paula...

o ist mit ihren Eltern in die Berge gefahren

o schwimmt gerne im Meer

o hat Angst vor Krebsen

o mag farbige Muscheln

Insgesamt sind 20 Rohwertpunkte erreichbar, die als Gesamtscore in der weiteren Analyse berücksichtigt werden. Jede richtig gelöste Aufgabe wird mit einem Punkt verrechnet. Wird eine Aufgabe nicht bearbeitet, die falsche Antwortalternative ausgewählt oder mehrere Antworten markiert, wird dies mit null Punkten bewertet.

6.2.3. Peabody Picture Vocabulary Test (PPVT-III)

Dieser Abschnitt wurde von Stefanie Dorn (2011) verfasst.

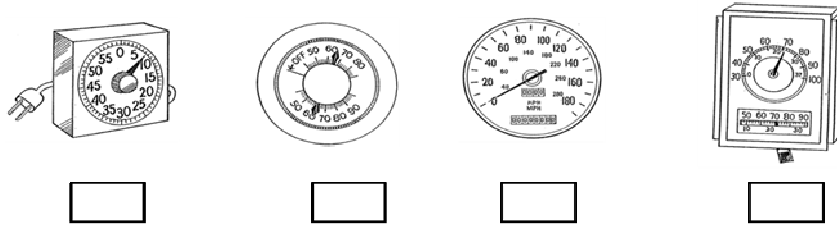
Der PPVT-III (Dunn & Dunn, 1997) ist die dritte Auflage eines Vokabeltests, der das rezeptive (Sprach-) Vokabular, den passiven Wortschatz, von Personen im Alter von 2;6 bis 90 Jahren erfasst. Der PPVT-III liegt für unterschiedliche Sprachen vor. Die insgesamt 204 Items sind in 17 Sets zu je 12 Items mit ansteigender Schwierigkeit angeordnet. Die Aufgabenstellung sieht vor, dasjenige Bild aus vier vorgegebenen Möglichkeiten zu wählen, das die Bedeutung des mündlich vorgetragenen Stimuluswortes am besten repräsentiert. Die Autoren geben für jede Altersgruppe ein Einstiegsset an, von dem ausgehend die Items der Reihe nach bearbeitet werden. Das Ende der Testung ist erreicht, wenn in einem Set acht oder mehr Fehler gemacht werden (bzw. das letzte Item, Nr. 204, erreicht ist), was im Durchschnitt nach fünf Sets (60 Items) der Fall ist (Dunn & Dunn, 1997).

Der PPVT-III ist ein sehr etabliertes und häufig verwendetes Verfahren, das auch den psychometrischen Ansprüchen genügt. Die Reliabilitäten für den Altersbereich der untersuchten Stichprobe liegen über $r = .90$ und die Ergebnisse der Validierungsstudien sind ebenfalls als ausreichend zu betrachten (z.B. WISC-III: $r = .82$ bis $r = .91$) (Dunn & Dunn, 1997).

Im Rahmen dieses Forschungsprojektes wurde der PPVT-III zur Erfassung der (rezeptiven) lexikalischen Fähigkeiten (vgl. auch Braze et al., 2007; Protopapas et al., 2007) beziehungsweise als Index für den Wortschatz eingesetzt.

Aus Gründen der Ökonomie, Zumutbarkeit und Vergleichbarkeit von Testergebnissen wurde die ursprüngliche Form des PPVT-III als Gruppentestung adaptiert (Anhang D). Zu diesem Zweck bekamen die SchülerInnen einen Antwortbogen vorgelegt, auf dem die Lösung und drei Distraktoren abgebildet waren (analog zur Originalfassung von Dunn & Dunn, 1997). Zur besseren Orientierung wurde die Bilderreihe zusätzlich als Power-Point-Präsentation im Klassenzimmer an eine Wand projiziert. Es wurde das Zielwort laut vorgelesen und die Kinder hatten 15 Sekunden lang Zeit (in Anlehnung an die Originalfassung von Dunn & Dunn, 1997) die richtige Lösung zu markieren.

Beispiel: „Schaltuhr“ (wurde mündlich vorgelesen)



Da im Rahmen dieses Projektes auch die Testung von Kindern der zweiten (siehe Klausecker, 2011) und dritten (siehe Huemer, 2011) Schulstufe vorgesehen war, wurde als Einstieg Set 5 (Altersbereich 6-7 Jahre) gewählt. Des Weiteren wurden Set 7 (Altersbereich 8-9), 8 (Altersbereich 10-11), 10 (Altersbereich 12-16) und 12 (Altersbereich 12-16) vorgegeben, sodass die Anzahl der durchschnittlich 60 Items zur Leistungserfassung realisiert werden konnte. Für die statistische Auswertung werden die Rohwerte (Anzahl korrekter Items) der TeilnehmerInnen verwendet, wobei jede richtige Antwort mit einem Punkt verrechnet wird.

6.2.4. Wortlesetest

Dieser Abschnitt wurde von Stefanie Dorn (2011) verfasst.

Der Wortlesetest von Schabmann und Kollegen (2009) ist ein effizientes Verfahren um die mündliche Lesefähigkeit zu überprüfen. Kognitionspsychologische Befunde der Leseforschung sehen die Wortleseleistung als einen wesentlichen Faktor der Lesefertigkeiten an, da Texte im Allgemeinen dann gut elaboriert werden können, wenn Wörter schnell und fehlerfrei identifiziert werden (Landerl & Willburger, 2009).

Im Sinne der Zwei-Wege-Theorie (vgl. Abschnitt 3.1.3.) kommt es durch den direkten lexikalischen Zugriff auf das mentale Lexikon zum schnellen und effizienten Worterkennen und gleichzeitig können Wörter/Buchstabenfolgen über den indirekten Weg rekodiert werden. Es wird davon ausgegangen, dass bei geübten LeserInnen die Worterkennung hochautomatisiert abläuft. Somit erscheint die Leseflüssigkeit (fluency) als geeigneter Indikator für die Worterkennung/Wortleseleistung. Die Lesegenauigkeit (accuracy) ist in Schriftsprachen mit hoher Graphem-Phonem-Korrespondenz bereits früh in der

Leseentwicklung ausgeprägt, sodass nahezu jedes Wort zuverlässig lautierend erlesen werden kann (Landerl & Willburger, 2009). Durch die Vorgabe von (Pseudo-)Wörtern ohne jeglichen semantischen Kontext wird zum einen die Fähigkeit der ProbandInnen zum schnellen, direkten Worterkennen überprüft, zum anderen wird die genaue und flüssige Wortrekodierung erfasst.

Durch die eindeutige Aufgabenstellung und die daraus resultierenden Leistungsanforderungen ist der Test als reliabel und valide einzuschätzen. Die vorgegebenen Instruktionen und die einfache Auswertung sichern darüber hinaus die Objektivität. In dieser Untersuchung wird der Wortlesetest als Indikator für die Leseflüssigkeit und Lesegenauigkeit herangezogen, wobei diese gemeinsam die Variable Worterkennen abbilden.

Den Kindern wurden im Einzelsetting 90 Wörter am PC dargeboten, die sie so schnell wie möglich und gleichzeitig fehlerfrei laut vorlesen sollten. Die Darbietungszeit der einzeln dargebotenen Wörter konnte von den Kindern durch Tastendruck (Erscheinen des neuen Wortes) selbst bestimmt werden. Die Wörter sind in sechs Blöcke zu je 15 Wörtern aufgeteilt und nach jedem absolvierten Wortblock wurde eine kurze Pause eingelegt. Die ersten 60 „realen“ Wörter steigen in ihrem Schwierigkeitsgrad kontinuierlich an und lassen sich anhand der Merkmale „häufig/selten vorkommend“ und „kurz/lang beziehungsweise einsilbig/dreisilbig“ unterscheiden. Die 30 Pseudowörter sind gemäß ihrer Wortlänge differenziert.

Beispiele (entnommen aus Form A, siehe Anhang D):

- Block 1-kurze, häufige Wörter: Post, Bad, Ziel,...
- Block 2-kurze, seltene Wörter: Schmutz, Spatz, Knie...
- Block 3-lange, häufige Wörter: Erziehung, Beamtin, Kartoffel...
- Block 4-lange, seltene Wörter: Omnibus, Heiterkeit, Kamerad...
- Block 5-kurze Pseudowörter: Faka, Breigt, Frilp, ...
- Block 6-lange Pseudowörter: Verbalut, Mandriche, Schelperta,...

Für die Auswertung wird ein Gesamtscore der Lesezeit (Fluency) und ein Gesamtscore der Lesefehler (Accuracy) über alle Wörter hinweg berechnet.

6.3. Untersuchungsdesign und statistische Analyse

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Bei der Datenanalyse der einmalig durchgeführten Gesamterhebung kommen zwei Untersuchungsstrategien zum Einsatz. Erstens wird ein deskriptivstatistisches Vorgehen zur Beschreibung der Stichprobe und der Testergebnisse gewählt. Zweitens eine prüfende Strategie zur Untersuchung der theoretisch postulierten Modellannahmen. Hierbei werden die Kausalzusammenhänge zwischen den latenten Variablen simultan mittels Strukturgleichungsmodellen (Structural Equation Models, SEM) überprüft. An dieser Stelle werden lediglich die Daten der vierten Klasse berücksichtigt. Für die Ergebnisse der zweiten und dritten Schulstufe sei auf Klausecker (2011) und Huemer (2011) verweisen.

Als manifeste Indikatoren für die latenten Variablen *Leseverständnis*, *Wortschatz* und *Hörverstehen* werden die Testwerte der entsprechenden Testverfahren verwendet (vgl. Abschnitt 6.2.). In der endogenen Variable *Leseverständnis* ist zusätzlich die nicht erklärbare, individuelle Residualvarianz (als resLV bezeichnet) berücksichtigt. Die latente erklärende Variable *Worterkennen* wird durch die Indikatorvariablen *fluency* und *accuracy* operationalisiert, wobei das Messmodell mit zwei Fehlertermen (err1 und err2) behaftet ist, welche als Messfehler der beiden Indikatorvariablen interpretiert werden können.

In jedem der Strukturmodelle (Modell 1, 2, 3; siehe Abschnitt 7.2.) wird die kausale Beziehung der hypothetischen Konstrukte abgebildet, wobei die Pfeilrichtungen die postulierte Kausalität der Variablen anzeigen. Wird keine direkte (kausale) Beziehung angenommen, so ist dies durch Korrelationen (Doppelpfeil) spezifiziert.

Die simultane Parameterschätzung für jedes Modell erfolgt durch die Maximum-Likelihood (ML) Methode, anhand derer inferenzstatistisch (χ^2 -Test) überprüft wird, ob die modelltheoretische Varianz-Kovarianzmatrix eine hinreichend gute Reproduktion der

empirischen Varianz-Kovarianzmatrix darstellt (Backhaus, Erichson, Plinke & Weiber, 2008; Buch, 2007; Reiner, 2009).

Zur Beurteilung der Anpassungsgüte („Fit“) des Gesamtmodells liegen zahlreiche statistische Indizes vor, wobei hier überblicksartig auf die wesentlichen eingegangen werden soll. Sofern die nachfolgenden Angaben nicht anderweitig gekennzeichnet sind, entstammen sie den Arbeiten von Buch (2007), Byrne (2001), Schermelleh-Engel, Moosbrugger & Müller (2003) und Reiner (2009). Der Chi-Quadrat (χ^2)-Test überprüft die Nullhypothese, dass die empirische Varianz-Kovarianzmatrix der modelltheoretischen entspricht. Liefert dementsprechend die χ^2 -Teststatistik ein nicht signifikantes Ergebnis, kann davon ausgegangen werden, dass sich die Matrizen entsprechen und das Modell einen guten Fit aufweist (Hoyle, 1995). Die Freiheitsgrade (df) werden im „normed chi-square“-Index (CMIN/df) berücksichtigt, indem der χ^2 -Wert durch diese dividiert wird.

Der Root-Mean-Square-Error of Approximation (RMSEA) basiert auf der Schätzung des Minimums der Diskrepanzfunktion (zwischen modelltheoretischer und empirischer Varianz-Kovarianzmatrix) und berücksichtigt zusätzlich die Modellkomplexität. Modelle mit Werten kleiner als .05 gelten als gut an die vorliegenden Daten angepasst. Wird für den RMSEA ein kleines Konfidenzintervall angegeben, so deutet dies auf die Genauigkeit dieses Wertes bezüglich seiner Aussagekraft über die Modellanpassung hin. Für den RMSEA wird ein p-Wert, als PCLOSE bezeichnet, berechnet. Nimmt der PCLOSE Werte größer als .50 an, gilt der RMSEA als verlässlicher Index für die vorliegende Stichprobe.

Der Goodness-of-Fit Index (GFI), der den relativen Anteil der erklärten Varianzen und Kovarianzen der Grundgesamtheit durch das Modell erfasst, kann analog dem Bestimmtheitsmaß der Regressionsanalyse interpretiert werden. Werte größer oder gleich .90 spiegeln eine gute Modellanpassung wider. Die Freiheitsgrade des Modells, und somit seine Komplexität, werden im Adjusted Goodness-of-fit Index (AGFI) berücksichtigt. Auch dieser Wert kann im Bereich von 0 bis 1 liegen, wobei höhere Werte für einen besseren Modellfit sprechen und dieser mindestens .90 erreichen sollte.

Mithilfe des Normed Fit Index (NFI) wird das aktuelle Modell mit dem schlecht möglichsten Modell (alle Variablen unkorreliert) und dem bestmöglichen Modell (alle möglichen

Parameterschätzungen enthalten) verglichen. Je ähnlicher das vorliegende Modell dem saturierten Modell ist, umso eher nimmt der NFI einen Wert von 1 an. Werte unter .90 deuten auf eine schlechte Anpassung des Modells hin. Im Comparative Fit Index (CFI) wird das Problem der Stichprobenabhängigkeit des NFI durch Berücksichtigung der Freiheitsgrade zu lösen versucht. Werte von mindestens .90 zeugen von einem zufriedenstellenden Modellfit.

An die Beurteilung des Gesamtmodells anschließend werden die einzelnen Modellteile sukzessive in ihrem Erklärungswert für die Varianzaufklärung der endogenen Variable analysiert. Die Differenz der beobachteten und geschätzten Werte (Residuen) sollte für jeden Parameter möglichst klein sein, allerdings wird eher die Betrachtung der standardisierten Residuen (Residuum dividiert durch den geschätzten Standardfehler) empfohlen, da diese unabhängig von der Skalierung der Variablen und somit besser interpretierbar sind. Werte größer als 2,58 sind als kritisch zu betrachten, da die Wahrscheinlichkeit für dieses Ergebnis bei Gültigkeit der Nullhypothese (Differenz gleich null) weniger als ein Prozent betragen würde.

Ob ein Parameter einen signifikanten Beitrag zur Modellstruktur liefert, kann anhand der Critical Ratio (C.R.) abgeschätzt werden. Der Wert der C.R. ist t-verteilt, sodass absolute Werte größer als 1,96 darauf hinweisen, dass der geschätzte Parameter sich signifikant von 0 unterscheidet (mit einer Irrtumswahrscheinlichkeit von 5%).

In jedem Modell wird darüber hinaus mittels Satorra Bentler-Test überprüft, ob es zu signifikanten Änderungen durch die Festsetzung/Einsparung von Korrelationen kommt. Deuten die Ergebnisse der „nested model comparison“ auf ein nicht signifikantes Ergebnis hin, so können diese aus Gründen der Sparsamkeit beziehungsweise Reduktion auf das Notwendigste („parsimony“; vgl. Kline, 1998) vernachlässigt werden. Der Modellfit verschlechtert sich bei Restringierung der Korrelationen nicht. In analoger Weise wird beim Vergleich der drei postulierten Modelle vorgefahren. Es wird überprüft, ob sich statistisch signifikante Unterschiede im Erklärungswert der Modelle bei Einsparung der angenommenen direkten Pfade zwischen Worterkennen und Wortschatz ergeben. Die Stärke aller direkten Effekte auf die latente Variable Leseverständnis kann anhand der quadrierten multiplen Korrelationen bewertet werden. Diese geben den durch das Modell aufgeklärten Varianzanteil an.

Die statistische Analyse wird mit den Statistikprogrammen SPSS 19.0 und AMOS 18 durchgeführt. Die Beantwortung der Fragestellungen erfolgt mit einer festgesetzten Irrtumswahrscheinlichkeit von 5% ($\alpha = .05$) anhand der statistischen Kennwerte der Strukturgleichungsmodelle.

7. Ergebnisdarstellung

Das Kapitel 7 wurde von Stefanie Dorn (2011) verfasst.

Die Ergebnisdarstellung umfasst eine allgemeine deskriptive Beschreibung der Stichprobe und der Testergebnisse der SchülerInnen der vierten Schulstufe. Im Anschluss daran werden die Daten dieser Stichprobe anhand von Strukturgleichungsmodellen analysiert. Die Ergebnisdarstellung der SchülerInnen der zweiten Schulstufe findet sich in der zeitgleich entstandenen Parallelarbeit von Klausecker (2011), die der SchülerInnen der dritten Schulstufe in der zeitgleich entstandenen Parallelarbeit von Huemer (2011).

7.1. Deskriptive Ergebnisdarstellung

7.1.1. Stichprobe

An der Untersuchung im Zeitraum von Dezember 2010 bis Februar 2011 nahmen insgesamt 362 VolksschülerInnen der Klassenstufe zwei, drei und vier aus dem Bezirk Vöcklabruck (Oberösterreich) teil. Die Daten wurden vier Volksschulen erhoben, wobei in der vorliegenden Diplomarbeit die Daten der 113 Kinder aus der vierten Schulstufe verwendet werden.

42,5% der Kinder (48) wurden in den ersten beiden Schulstunden getestet. In der dritten und vierten Schulstunde (nach der Pause zwischen zweiter und dritter Schulstunde) wurden 16,8% der Testungen durchgeführt. Bei 15 der teilnehmenden Kinder (13,3%) wurde die Testung durch die Pause unterbrochen. In diesem Fall wurde in der zweiten Schulstunde die Gruppentestung durchgeführt und in der dritten Stunde der Lesetest vorgegeben. 14,2% der

Kinder absolvierten die Testung in der vierten und fünften Schulstunde und bei 15 Kindern wurde die Testung (der Lesetest) an einem anderen Tag weitergeführt.

An der Untersuchung nahmen 51 Knaben (45,1%) und 62 Mädchen (54,9%) im Alter von 8 bis 11 Jahren ($MW = 9,65$, $Median = 10$, $SD = 0,55$) teil. Der Großteil der SchülerInnen war zum Testzeitpunkt 9 Jahre (36,3%) oder 10 Jahre (60,2%) alt, lediglich ein Schüler war jünger (8 Jahre) und drei SchülerInnen älter (11 Jahre). Die Altersverteilung getrennt für beide Geschlechter ist in Abbildung 2 ersichtlich. 37,3% (19) der teilnehmenden Knaben waren 9 Jahre alt sowie 35,5% der Mädchen. Mehr als die Hälfte der untersuchten Stichprobe war 10 Jahre alt (56,9% der Knaben, 62,9% der Mädchen).

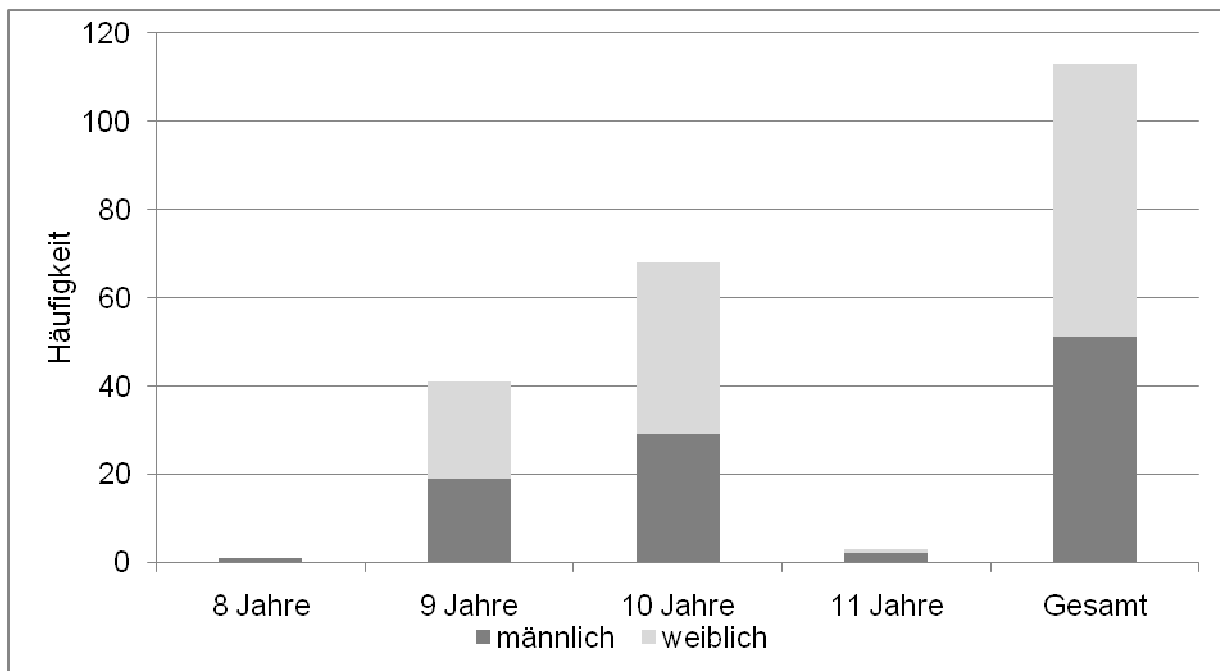


Abbildung 2: Geschlechtsspezifische Darstellung der Altersverteilung (N=113)

Die Stichprobe besteht zum überwiegenden Anteil aus deutschsprachigen Kindern (94,7%), nur 6 SchülerInnen gaben eine andere Muttersprache als Deutsch an. Ein Schüler erhält sonderpädagogische Förderung, eine diagnostizierte Legasthenie lag in keinem der Fälle vor.

In der Datenanalyse wurden alle 113 Viertklässler berücksichtigt, da kein Kind in sämtlichen Testverfahren durchwegs sehr schlechte Leistungen erzielte. Lediglich in einzelnen Tests wurden von verschiedenen SchülerInnen unterdurchschnittliche Leistungen erbracht. Diese wurden jedoch im Hinblick auf eine ausreichend große Stichprobe zur Berechnung der Strukturgleichungsmodelle nicht aus dem Datensatz entfernt.

7.1.2. Ergebnisse der Testverfahren

Die Ergebnisse des Subtests „Textverständnis“ aus dem ELFE 1-6 (Lenhard & Schneider, 2006), des Subtests „Hörverstehen“ aus dem KnuspeL (Marx, 1998), des Wortschatztests (PPVT-III, Dunn & Dunn, 1997) und des Wortlesetests (Schabmann et al., 2009) werden anhand ihrer deskriptiven Kennzahlen näher erläutert.

Insgesamt wurden die Testrohre von 113 SchülerInnen für die jeweilige Variable analysiert und in Tabelle 1 anhand des arithmetischen Mittels (Mittelwert, MW), der Standardabweichung (SD), des Minimums (Min) und des Maximums (Max) dargestellt. Der Median für jeden Testwert wird nicht weiter berücksichtigt, da sich dieser nur minimal vom Mittelwert unterscheidet. Auf eine geschlechtsspezifische Darstellung der Daten wird verzichtet, da diese erstens im Rahmen der weiteren Verfahren irrelevant sind und zweitens die Unterschiede im Mittelwert nur für das Hörverstehen statistisch signifikant ($p = .040$, $\alpha = .05$) ausfielen.

Die Leistung im *Worterkennen* setzt sich aus der Fähigkeit zum flüssigen (fluency) und fehlerfreien (accuracy) Lesen von Wörtern und Pseudowörtern zusammen. Die Variable wird durch die Gesamtfehleranzahl (accuracy) und die benötigte Zeit in Sekunden (fluency) im Wortlesetest über alle 90 Wörter hinweg berechnet. Im Mittel wurden 5,66 Fehler (SD= 4,49) beim lauten Vorlesen der Wörter gemacht. 7 Kinder konnten die Wörter fehlerfrei artikulieren, bei einem Kind waren 28 Lesefehler zu verzeichnen. Am häufigsten wurden drei Wörter fehlerhaft vorgelesen (15 Kinder, 13,3%). Circa zwei Drittel der SchülerInnen benötigten zum Lesen der 90 Wörter weniger als 141 Sekunden. Das schnellste Kind brauchte zum Vorlesen etwas mehr als eine Minute (62,5 Sekunden), das Langsamste benötigte etwas mehr als vier Minuten (242,16 Sekunden). Im Mittel betrug die Lesezeit für alle 6 Wortblöcke zwei Minuten und 14 Sekunden, wobei die Standardabweichung eine halbe Minute beträgt

(SD = 30,32). Angaben zur Normstichprobe können aufgrund fehlender Daten nicht gemacht werden.

Der *Wortschatz* der vorliegenden Stichprobe stellt sich als sehr homogen dar. Von den insgesamt 60 erreichbaren Punkten im adaptierten Wortschatztest (PPVT-III) wurden durchschnittlich 50,81 Punkte (SD = 3,25) erzielt, wobei 81,3% (92) der Kinder einen Gesamtscore zwischen 48 und 55 Rohwertpunkten erreichten. Ein Kind konnte nur 42 der Aufgaben richtig lösen und zwei SchülerInnen erzielten das Maximum in dieser Stichprobe mit 57 korrekt gelösten Items. Aufgrund der angepassten Vorgabe des Verfahrens ist auch hier ein Vergleich mit der Normstichprobe nicht möglich.

Die Leistungen im *Hörverstehen* sind ebenfalls als sehr homogen zu bezeichnen. Im Hörverstehenstest (Knuspel-L; Marx, 1998) erzielten 73,46% der untersuchten Stichprobe einen Wert zwischen 30 und 34 Punkten. Die insgesamt 34 Aufgaben wurden von drei SchülerInnen korrekt bearbeitet, die schlechteste Testleistung (19 Punkte) wurde in dieser Stichprobe einmal erzielt. Die Standardabweichung der Testwerte (SD) vom arithmetischen Mittel (MW = 30,36) beträgt 2,49. Für die Normstichprobe (n=302) wird das arithmetische Mittel mit 30,02 angegeben, die Standardabweichung mit einem Wert von 3,09. Bei getrennter Betrachtung der Testwerte nach Wissens- und Ausführungsaspekten der Items (maximal je 17 Punkte), ergibt sich für das arithmetische Mittel des Wissensaspektes ein Wert von 14,97 (Normstichprobe MW = 14,78) und eine Standardabweichung von 1,37 (Normstichprobe SD = 1,45). Der Mittelwert des Ausführungsaspekts der Items beträgt 15,39 (Normstichprobe MW = 15,24) und für die Standardabweichung ergibt sich ein Wert von 1,58 (Normstichprobe SD = 1,64).

Die Bestleistung von 20 zu erreichenden Punkten in der Variable *Textverständnis*, erhoben mittels Textverständnistest aus dem ELFE 1-6 (Lenhard & Schneider, 2006), konnte von sechs Kindern (5,3%) erzielten werden. Ein Schüler konnte nur ein Item richtig beantworten. Für die Normstichprobe (n = 275) finden sich im Manual folgende Angaben für Mittelwert (MW) und Standardabweichung (SD): MW = 15,05 und SD = 4,05. In der vorliegenden Stichprobe der Viertklässler wurden im Mittel 14,70 Punkte erzielt (SD = 3,65), wobei 54% der SchülerInnen einen Testwert von weniger als 16 Rohwertpunkten erreichten.

Tabelle 1: Deskriptive Beschreibung der Testergebnisse für die jeweilige Variable der Stichprobe (N = 113)

Variable	MW	SD	Min	Max
Accuracy (Fehler)	5,66	4,49	0	28
Fluency (Sekunden)	133,65	30,32	62,50	242,16
Wortschatz	50,81	3,25	42	57
Hörverstehen	30,36	2,49	19	34
Textverständnis	14,70	3,65	1	20

7.2. Ergebnisdarstellung der Strukturgleichungsmodelle

Aufgrund theoretischer Überlegungen wurde bereits in Kapitel 5.2. (Fragestellungen) ein hypothetisches Modell konstruiert, von dem ausgehend drei Modelle abgeleitet wurden, die sich in ihren direkten Pfaden unterscheiden. Für diese Modelle wurden zunächst mögliche Korrelationen zwischen den Prädiktorvariablen *Worterkennen*, *Wortschatz* und *Hörverstehen* angenommen. Wie in Tabelle 2 ersichtlich erwies sich der Zusammenhang zwischen den Variablen *Wortschatz* und *Hörverstehen* für alle drei Modelle als substantiell. Diese signifikante Korrelation muss folglich in der weiteren Ergebnisdarstellung berücksichtigt werden, um nicht eine Verschlechterung des Modellfits zu riskieren. Aus Gründen der Sparsamkeit werden die nicht signifikanten Korrelationen nicht weiter analysiert, da bei deren Auslassung keine signifikante Verschlechterung des Modellfits zu erwarten ist.

Tabelle 2: Korrelationen im Modellvergleich

Modell	Korrelation	p
Modell 1	Worterkennen – Wortschatz	.524
Modell 1	Worterkennen – Hörverstehen	.575
Modell 1	Wortschatz – Hörverstehen	.000
Modell 2	Worterkennen – Hörverstehen	.341
Modell 2	Wortschatz – Hörverstehen	.000
Modell 3	Worterkennen – Hörverstehen	.575
Modell 3	Wortschatz- Hörverstehen	.000

In den nachfolgenden Abschnitten werden die Resultate der Parameterschätzungen zu den drei Modellen angegeben, wobei die geforderten Verfahrensvoraussetzungen als erfüllt betrachtet werden. Wie in Abschnitt 6.3. (Untersuchungsdesign und statistische Analyse) erläutert, wird für jedes Modell ein Chi-Quadrat (χ^2)-Test durchgeführt. Liefert die χ^2 -Teststatistik ein nicht signifikantes Ergebnis, kann von einem guten Modellfit ausgegangen werden. Die Anpassungsindices GFI, AGFI, NFI und CFI sollten in gut angepassten Modellen einen Wert größer oder gleich .90 erreichen. Vom Wert des RMSEA, der in einem kleinen Konfidenzintervall mit einem PCLOSE-Wert größer als .50 liegen sollte, ist für eine gute Modellanpassung zu fordern, dass er keine Werte größer .05 annimmt. Die einzelnen Teilstrukturen der Modelle werden anhand ihrer standardisierten Residuen (Werte kleiner als 2,58 sind akzeptabel) und der C.R. (Werte größer als 1,96 sind auf einem 5%igen Signifikanzniveau statistisch signifikant) beurteilt. Die standardisierten Regressionsgewichte (Koeffizienten p_{ij}) geben die Höhe/Stärke des direkten Einflusses/Effektes der einzelnen Prädiktoren auf die gemessene Leseverständnisleistung wider. Die Varianzen der latenten Variablen sind auf 1 fixiert, sodass es sich um standardisierte Lösungen handelt. Die Stärke der Effekte der einzelnen Variablen ist somit vergleichbar. Die Ergebnisse der einzelnen Modelle (AMOS-Output) sind in Anhang E zu finden.

7.2.1. Modell 1

Die in Modell 1 postulierten Annahmen spiegeln die in der Realität erhobenen Daten gut wider (χ^2 -Test: $\chi^2 = 2,827$, $df = 4$, $p = .587$). Der Varianzanteil, der durch das Modell reproduziert werden kann, ist als hoch zu bewerten ($GFI = .990$). Der Wert des NFI (.973) weist darauf hin, dass das Modell 1 sehr nahe am saturierten Modell liegt und auch der CFI (1.000) und der AGFI (.963) lassen auf eine gute Gesamtanpassung des Modells schließen. Der Wert des RMSEA (.000) liegt in einem kleinen Konfidenzintervall [.000; .122] und kann als verlässlicher Index für die untersuchte Population gelten ($PCLOSE = .705$). Die Werte der standardisierten Residuen liegen unter dem kritischen Wert von 2,58.

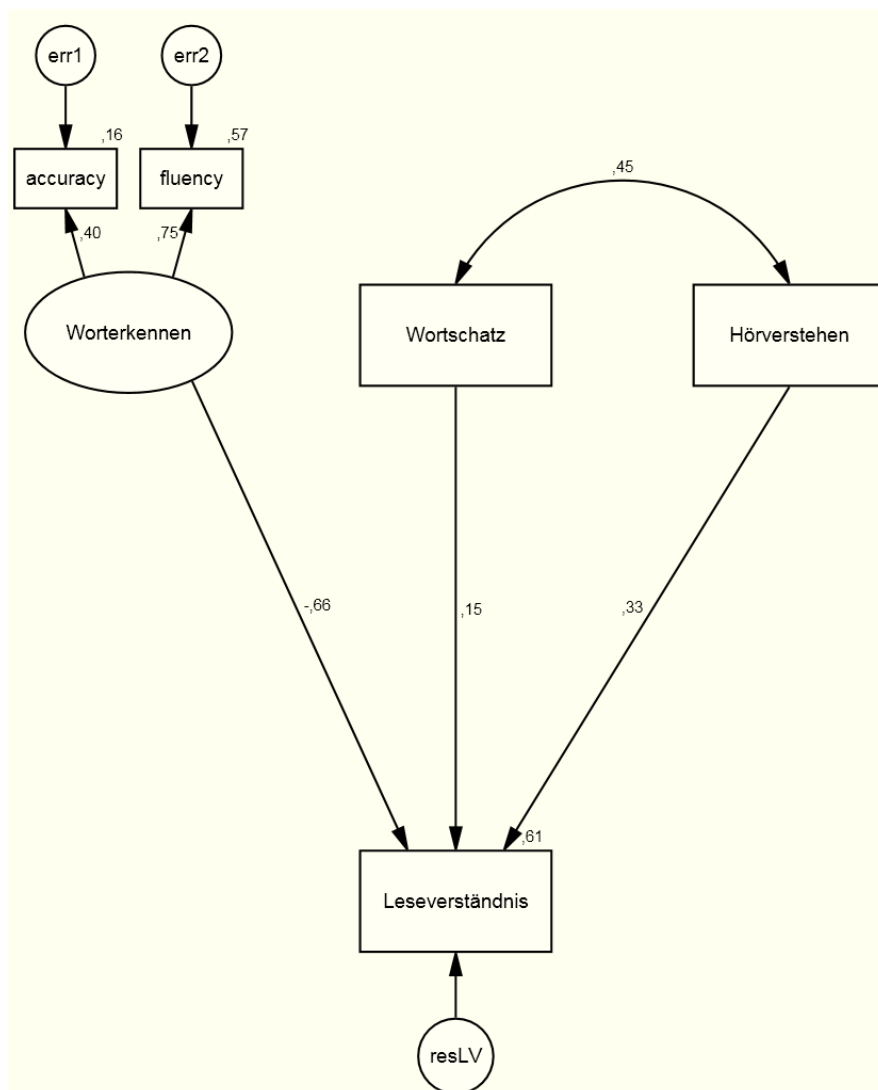


Abbildung 3: Modell 1 mit standardisierten Lösungen

Wie in Abbildung 3 erkennbar ist, können 61% der Varianz im *Leseverständnis* durch die Prädiktoren erklärt werden. Der stärkste Prädiktor des *Leseverständnisses* ist das *Worterkennen* (C.R. = -3,212, $p = .001$) mit hohem Koeffizienten ($\beta_{ij} = -.66$). Der negative Wert resultiert aus der Skalierung der Variablen *accuracy* (Fehleranzahl) und *fluency* (Lesezeit). Das *Hörverstehen* leistet ebenfalls einen signifikanten Beitrag zum *Leseverständnis* (C.R. = 4,189, $p = .000$) mit mittelmäßigem Koeffizienten ($\beta_{ij} = .33$) und die substantielle Korrelation mit dem *Wortschatz* liegt im mittleren Bereich ($r = .45$). Der Koeffizient des *Wortschatzes* zeigt mit $\beta_{ij} = .15$ einen geringen Wert an. Obwohl sich der Einfluss des *Wortschatzes* auf das *Leseverständnis* als nicht signifikant erweist (C.R. = 1,891, $p = .059$), bleibt dieser als determinierende Variable im Modell enthalten (nähere Erläuterungen in Abschnitt 8. Diskussion). Das *Leseverständnis* kann als durch den Beitrag von *Worterkennen*, *Wortschatz* und *Hörverstehen* direkt beeinflusst gelten, zwischen den Variablen *Wortschatz* und *Hörverstehen* besteht ein positiver korrelativer Zusammenhang.

7.2.2. Modell 2

Modell 2 (siehe Abbildung 4) unterscheidet sich von Modell 1 durch die Annahme eines direkten Effektes vom *Wortschatz* auf das *Worterkennen*. Dieser erwies sich jedoch als sehr gering und statistisch nicht signifikant ($\beta_{ij} = .08$; C.R. = .643; $p = .520$).

Das *Worterkennen* weist mit einem hohen Koeffizienten von $\beta_{ij} = -.66$ den größten Einfluss auf das *Leseverständnis* auf, der Koeffizient des *Hörverstehens* liegt mit $\beta_{ij} = .34$ im mittleren Bereich. Der Koeffizient des *Wortschatzes* beträgt $\beta_{ij} = .17$. Die standardisierten Residuen liegen im akzeptablen Bereich und die Werte der Critical Ratio deuten darauf hin, dass die Prädiktorvariablen *Worterkennen* (C.R. = -3,198, $p = .001$), *Wortschatz* (C.R. = 1,970, $p = .049$) und *Hörverstehen* (C.R. = 4,195, $p = .000$) einen signifikanten Beitrag zur Modellstruktur liefern. Der Korrelationskoeffizient zwischen *Wortschatz* und *Hörverstehen* liegt bei $r = .45$. Durch die drei Prädiktoren werden 60% der Varianz des *Leseverständnisses* erklärt, die restlichen 40% sind durch Messfehler und andere Einflussfaktoren bedingt.

Insgesamt resultiert für das Modell mit $\chi^2 = 2,404$ und $df = 3$ eine nicht signifikante Modellabweichung ($p = .493$) von den erhobenen Daten. Die Indizes zur Beurteilung der Gesamtgüte des Modells weisen alle auf ein zufriedenstellendes Modell hin (GFI = .992;

AGFI = .958; NFI = .977; CFI = 1,000). Das Modell kann als gut an die Daten angepasst gelten (RMSEA= .000), und dieser Wert scheint präzise in seiner Aussagekraft über die Modellanpassung zu sein (Konfidenzintervall des RMSEA [.000; .147], PCLOSE = .605).

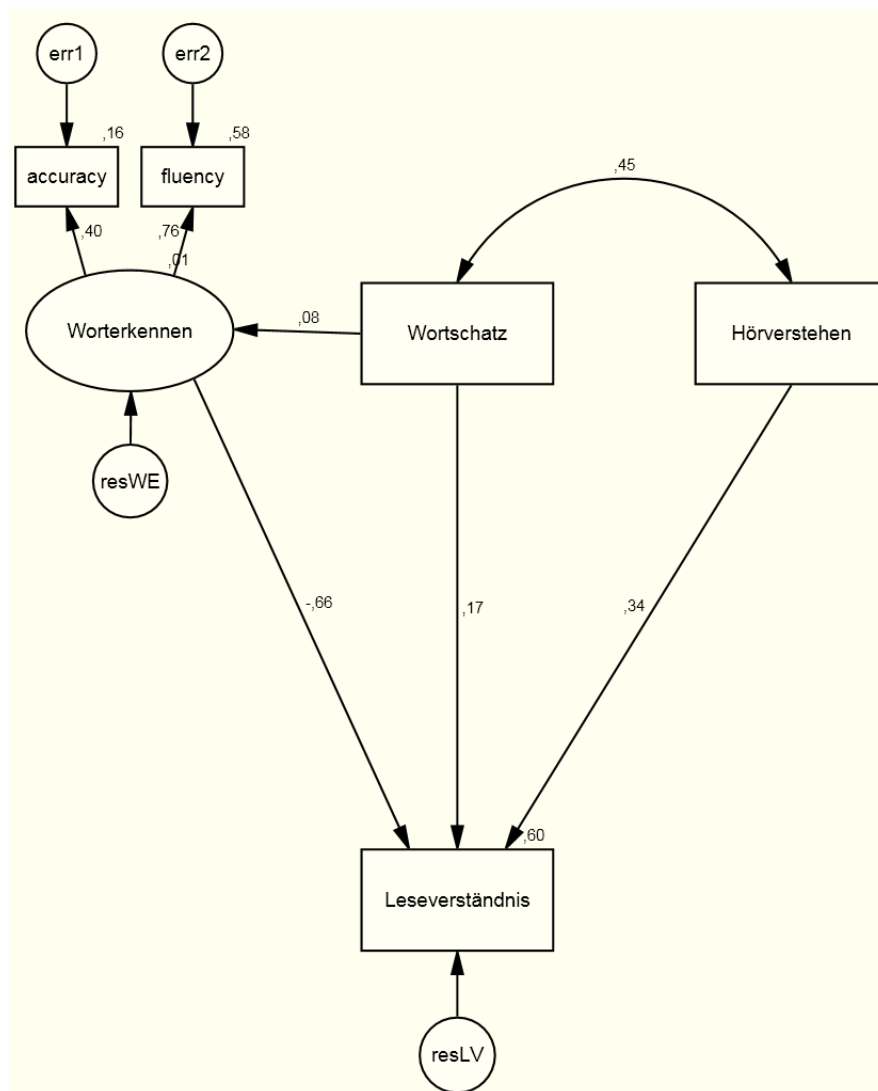


Abbildung 4: Modell 2 mit standardisierten Lösungen

7.2.3. Modell 3

Modell 3 (siehe Abbildung 5) prüft zusätzlich zu Modell 1 den direkten Effekt des *Worterkennens* auf den *Wortschatz*. Dieser Pfad liefert keinen signifikanten Erklärungswert für das Modell (C.R. = .973, $p = .330$) und sein Koeffizient ist niedrig ($p_{ij} = .11$).

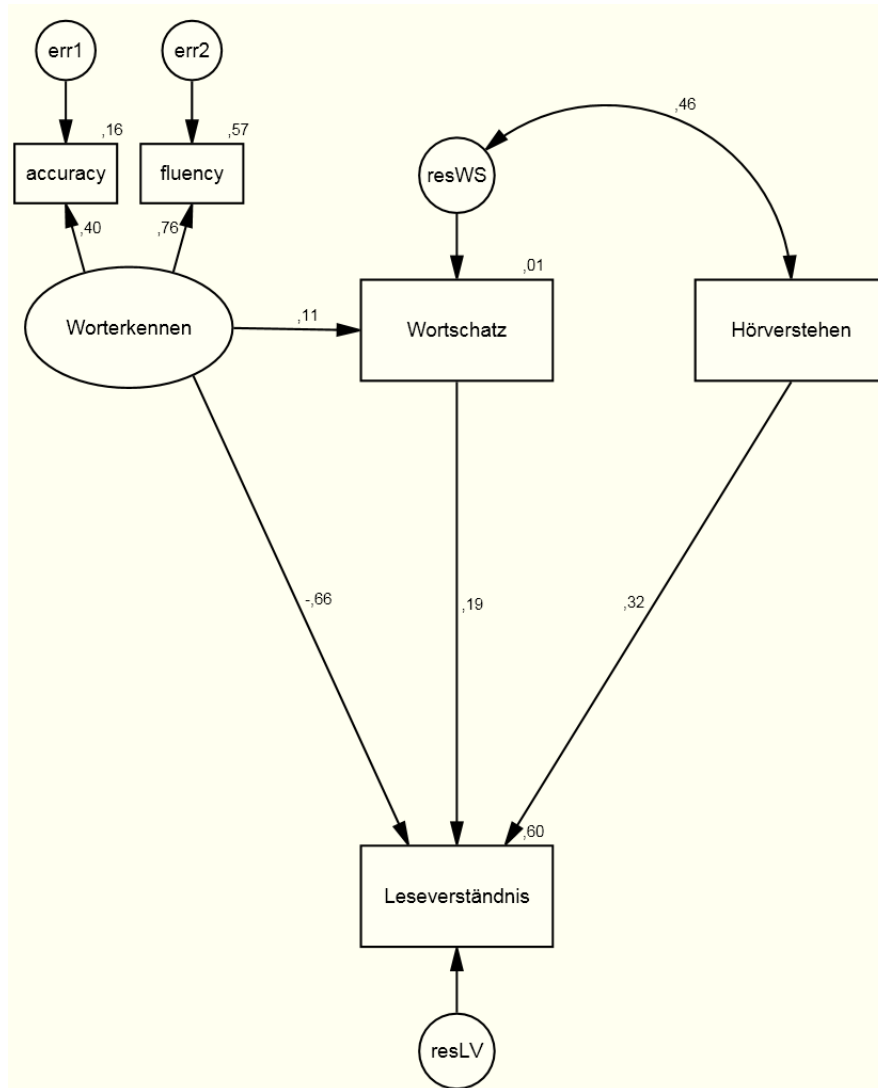


Abbildung 5: Modell 3 mit standardisierten Lösungen

Auch in Modell 3 ist das *Worterkennen* der stärkste Prädiktor mit hohem Koeffizienten ($p_{ij} = -.66$), der signifikant zur Varianzaufklärung des *Leseverständnisses* beiträgt (C.R. = -3,191, $p = .001$). Die Koeffizienten des *Hörverstehens* ($p_{ij} = .32$) und des *Wortschatzes* ($p_{ij} = .19$) liegen im niedrigen bis mittelmäßigen Bereich, wobei für beide ein statistisch signifikanter Einfluss nachgewiesen werden kann (Hörverstehen: C.R. = 3,854, $p = .000$; Wortschatz: C.R. = 2,049,

$p = .040$). *Wortschatz* und *Hörverstehen* korrelieren in einer Höhe von $r = .46$. Die Werte der standardisierten Residuen liegen unter dem kritischen Betrag von 2, 58. Die drei Prädiktoren können gemeinsam 60% der Varianz im *Leseverständnis* aufklären, der unerklärte Varianzanteil von 40% ist auf Messfehler und andere (nicht berücksichtigte) Variablen zurückzuführen.

Die empirische und modelltheoretische Varianz Kovarianzmatrix sind als annähernd gleich zu betrachten (χ^2 -Test: $\chi^2 = 1,813$, $df = 3$, $p = .612$), sodass von einem guten Modellfit ausgegangen werden kann. Das Modell gilt als gut an die Daten angepasst ($GFI = .994$) und ist in der Lage einen hohen Varianzanteil aufzuklären ($AGFI = .968$). Modell 3 liegt sehr nahe am saturierten Modell ($NFI = .983$) und auch der CFI ($CFI = 1,000$) deutet auf eine gute Gesamtanpassung des Modells hin. Die Diskrepanz zwischen modelltheoretischer und empirischer Varianz Kovarianzmatrix ist sehr gering ($RMSEA = .000$) und der Wert des RMSEA kann als sehr genau (Konfidenzintervall des RMSEA $[.000; .131]$) und verlässlich ($PCLOSE = .709$) bezüglich seiner Aussagekraft über die Modellanpassung erachtet werden.

7.2.4. Modellvergleich

Jedes der drei Modelle ist in der Lage die empirischen Daten hinreichend gut zu erklären und die globalen Anpassungsmaße deuten alle auf entsprechend guten Modellfit hin (siehe Tabelle 3). Die Varianzaufklärung der einzelnen Modelle bewegt sich um 60%.

Tabelle 3: Vergleich der drei Strukturgleichungsmodelle

Modell	χ^2	df	p	CFI	GFI	AGFI	RMSEA
Modell 1	2, 827	4	.587	1	.990	.963	.000
Modell 2	2,404	3	.493	1	.992	.958	.000
Modell 3	1,813	3	.612	1	.994	.968	.000

Die Modelle unterscheiden sich durch die zusätzliche Annahme eines direkten Pfades vom *Wortschatz* auf das *Worterkennen* (Modell 2) beziehungsweise vom *Worterkennen* auf den *Wortschatz* (Modell 3), wobei diesen jeweils ein geringer, statistisch nicht signifikanter Beitrag zur Modellerklärung zukommt. Um das geeignetste und gleichzeitig sparsamste

Modell zu identifizieren, wird statistisch überprüft, ob die Prädiktionskraft der Modelle durch die Entfernung der Pfade eingeschränkt wird.

Zunächst wurde das Modell 2 mittels „nested model comparison“ dem Modell 1 gegenübergestellt. Indem der direkte Pfad vom *Wortschatz* auf das *Worterkennen* mit 0 fixiert wird, kann überprüft werden, ob dieser einen signifikanten Beitrag zur Modellbildung liefert (siehe Abbildung 6).

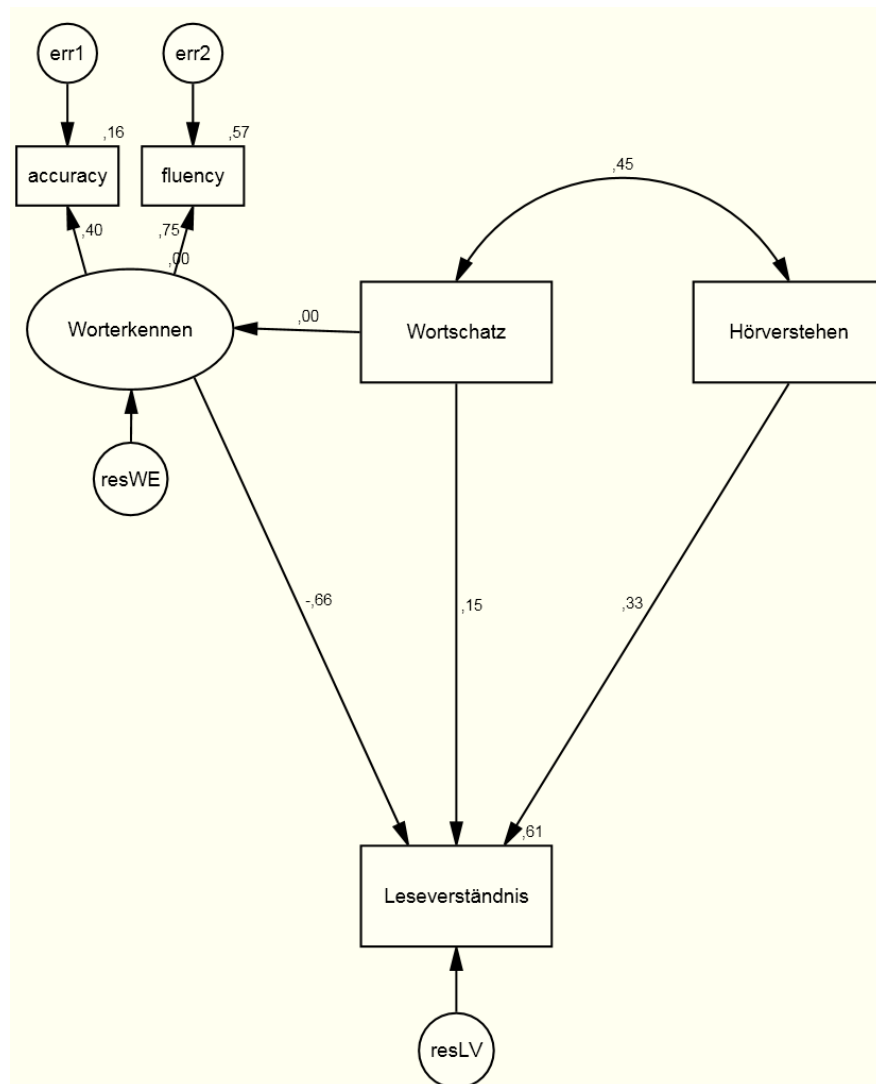


Abbildung 6: Nested model comparison für Modell 1 und Modell 2 mit standardisierten Lösungen

Das nicht signifikante Ergebnis ($p = .515$) des Satorra Bentler-Tests lässt darauf schließen, dass dieser Pfad keinen wesentlichen Beitrag zur Varianzaufklärung liefert und durch seine Entfernung aus dem Modell keine signifikante Verschlechterung des Modellfits resultiert. Aus Gründen der Sparsamkeit wird demzufolge Modell 1 angenommen.

In analoger Weise wurde mit Modell 3 verfahren (siehe Anhang E). Auch hier lieferte der Satorra Bentler-Test ein nicht signifikantes Ergebnis ($p = .314$). Die Unterschiede zwischen den Modellen sind statistisch nicht bedeutsam, sodass für die vorliegenden Daten die Gültigkeit von Modell 1 postuliert wird.

Schlussfolgernd lässt sich feststellen, dass in der vorliegenden Stichprobe das *Leseverständnis* direkt durch das *Worterkennen*, den *Wortschatz* und das *Hörverstehen* beeinflusst wird. Des Weiteren besteht eine substantielle mittlere Korrelation zwischen *Wortschatz* und *Hörverstehen* ($r = .45$). Den größten direkten Einfluss mit hohen Koeffizienten übt das *Worterkennen* aus. Der Einfluss des *Hörverstehens* und des *Wortschatzes* liegt im unteren bis mittleren Bereich (siehe Tabelle 4). Obwohl für den direkten Pfad zwischen *Wortschatz* und *Leseverständnis* keine statistische Signifikanz (bei $\alpha = .05$) nachgewiesen werden konnte, bleibt er aufgrund seiner praktischen Relevanz im Modell enthalten. Die direkten Effekte der erklärenden Variablen auf das *Leseverständnis* erweisen sich in ihrer Größenordnung als sehr robust gegenüber den Veränderungen der Spezifikationen in Modell 1 bis 3.

Tabelle 4: Koeffizienten, Critical Ratio (C.R.) und Signifikanz (p) der direkten Pfade im präferierten Modell 1

Pfade	β_{ij}	CR	p
Worterkennen → Leseverständnis	-.66	-3,212	.001
Wortschatz → Leseverständnis	.15	1,891	.059
Hörverstehen → Leseverständnis	.33	4,189	< .001

8. Diskussion

Der nachfolgende Abschnitt enthält die Interpretation der Ergebnisse für Kinder der vierten Schulstufe. Die Ergebnisse und Interpretationen bezüglich der Analyse der SchülerInnen der zweiten und dritten Schulstufe sind im entsprechenden Abschnitt bei Klausecker (2011) und Huemer (2011) zu finden. Abschnitt 8.2. beinhaltet eine kurze Darstellung der Ergebnisse des Gesamtprojektes. Abschließend werden kritische Anmerkungen zur durchgeführten Studie, sowie ein Ausblick auf zukünftige Replikationen und die Relevanz der vorliegenden Erhebung thematisiert.

8.1. Interpretation

Verfasst von Stefanie Dorn (2011).

Das angestrebte Ziel der vorliegenden Untersuchung bestand in der Integration des „simple view of reading“-Ansatzes (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000), namentlich der Faktoren *Worterkennen* und *Hörverstehen*, und des Wortschatzes, in Anlehnung an die „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002; Protopapas et al., 2007), in ein gemeinsames Modell zur Prädiktion des Leseverständnisses bei Kindern der vierten Volksschulklasse.

Die Ergebnisse der Strukturgleichungsmodelle legen nahe, dass das *Leseverständnis* direkt durch das *Worterkennen*, den *Wortschatz* und das *Hörverstehen* beeinflusst wird, wobei zusätzlich von einem mittleren positiven korrelativen Zusammenhang zwischen *Wortschatz* und *Hörverstehen* auszugehen ist. Der Effekt des *Worterkennens* liegt im mittleren bis hohen Bereich, der des *Hörverstehens* ist als klein bis mittelmäßig zu beurteilen und der *Wortschatz* scheint das *Leseverständnis* nur in geringem Ausmaß direkt zu beeinflussen. Die Annahmen, dass das *Worterkennen* zusätzlich einen direkten Effekt auf den *Wortschatz* ausübt (vgl. Protopapas et al., 2007) beziehungsweise dass der *Wortschatz* das *Worterkennen* direkt beeinflusst (vgl. Perfetti & Hart, 2002; Perfetti, 2007) konnten für die vorliegende Stichprobe nicht bestätigt werden.

Die Befunde können zunächst in Übereinstimmung mit den Annahmen des „simple view of reading“-Ansatzes gedeutet werden. Für das *Worterkennen* und das *Hörverstehen* konnte

jeweils ein unabhängiger Beitrag für das *Leseverständnis* nachgewiesen werden (vgl. auch Curtis, 1980), so dass diese beiden Komponenten als essentiell für das Leseverständnis zu betrachten sind (vgl. auch Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000; Rost & Hartmann, 1992).

Insbesondere Prozesse auf Wortebene scheinen hier entscheidenden Einfluss auf das Leseverständnis auszuüben. Das *Worterkennen*, das im Rahmen dieser Erhebung als durch die *Lesegenauigkeit* (*accuracy*) und *Leseflüssigkeit* (*fluency*) bedingt angenommen wird, trägt am meisten zur Aufklärung der Varianz im *Leseverständnis* bei. Demzufolge erleichtert eine schnelle und gleichzeitig sichere Wortidentifikation den Leseverständnisprozess. Kinder, die in der Lage sind schnell und fehlerfrei zu lesen, erzielen bessere Leistungen im Worterkennen, woraus in Folge höhere Verständnisseleistungen resultieren. Des Weiteren legen die Ergebnisse nahe, dass der Lesegeschwindigkeit gegenüber der Lesegenauigkeit größere Bedeutung beim Worterkennen zukommt. Dies kann als Beleg für eine fortschreitende Automatisierung des Worterkennungsprozess in der vorliegenden Stichprobe gewertet werden. Diese Ergebnisse und Schlussfolgerungen können durch verschiedene empirische Befunde und theoretischen Annahmen unterstützt werden. Zum einen ist aufgrund der hohen Konsistenz der Graphem-Phonem-Zuordnung in der deutschen Sprache bereits bei Kindern am Ende der ersten Klasse die Lesegenauigkeit sehr hoch (Klicpera & Gasteiger-Klicpera, 1993), sodass die Leseflüssigkeit beziehungsweise Geschwindigkeit als zentrales Kriterium für das Worterkennen gilt (Landerl & Willburger, 2009). Zum anderen kommt es durch die Automatisierung des lexikalischen und nichtlexikalischen Lesezugangs (Worterkennungsprozess) zur Steigerung der Lesegeschwindigkeit und einer Abnahme der Lesefehler (Klicpera et al., 2010). Kann auf Wörter respektive ihre lexikalischen Repräsentationen im Gedächtnis effizient (direkt und sicher) zugegriffen werden, dann stehen vermehrt kognitive Ressourcen für Verarbeitungs- und Verständnisprozesse auf höheren Ebenen zur Verfügung. Folglich sollte sich eine Zunahme der Lesegeschwindigkeit positiv auf das Leseverständnis auswirken (Cromley & Azevedo, 2007; Fuchs et al., 2001; Hosp & Fuchs, 2005). Der Einfluss des Dekodierens auf das Leseverständnis über die gesamte Grundschulzeit hinweg konnte unter anderem von Marx und Jungmann (2000) nachgewiesen werden. Die Autoren führen die zunehmende Leseverständnisseleistung auf die Automatisierung im Bereich des Dekodierens zurück. Auch sind die Ergebnisse mit der

„verbal efficiency theory“ von Perfetti (1985) kompatibel, in der davon ausgegangen wird, dass die Effizienz im Worterkennen den Erfolg des Leseprozesses bestimmt. Braze et al. (2007) konnten ebenfalls einen eigenen Beitrag des Dekodierens zum Leseverständnis feststellen.

Der signifikante Einfluss des *Hörverstehens* auf das *Leseverständnis* konnte für die Stichprobe ebenfalls nachgewiesen werden. Als Hörverstehen wird die Fähigkeit bezeichnet, sprachliche Informationen auf der Wortebene zur Ableitung von Satz- und Textinterpretationen über das Ohr zu nutzen (vgl. Marx & Jungmann, 2000). Das im Abschnitt 7.2.4. postulierte Modell des Leseverständnisses lässt die Schlussfolgerung zu, dass gut ausgeprägte Fähigkeiten im Hörverstehen zu besseren Leistungen im Leseverständnis führen. Dem Hörverstehen kommt demnach eine entscheidende Funktion im Leseverständnisprozess zu, was sich unter anderem bei Hagtvét (2003), Marx und Jungmann (2000) und Rost und Hartmann (1992) bestätigt findet.

Interessant in diesem Zusammenhang ist die Tatsache, dass das *Worterkennen*, im Vergleich zum *Hörverstehen*, einen doppelt so hohen Effekt auf das *Leseverständnis* ausübt. Marx und Jungmann (2000) gehen davon aus, dass insbesondere bei LeseanfängerInnen Schwierigkeiten im Worterkennen zu Problemen im Leseverständnis führen. Bei geübten LeserInnen entwickeln sich, durch die Auseinandersetzung mit der Schriftsprache, die Fertigkeiten zum Worterkennen kontinuierlich weiter, sodass das Hörverstehen die obere Leistungsgrenze im Leseverstehen bildet. Aufgrund der vorliegenden Daten würde sich aus der Stichprobe schließen lassen, dass das Worterkennen hauptsächlich durch die Lesegeschwindigkeit bedingt ist und somit die Automatisierung im Bereich Worterkennen repräsentiert wird. Die effiziente Worterkennung, in diesem Sinne als abhängig von der Erfahrung und Fähigkeit im Umgang mit Sprache zu interpretieren, bildet sich erst im Laufe der schulischen Entwicklung heraus und ist nach vier Jahren keineswegs abgeschlossen. Die Ergebnisse deuten lediglich darauf hin, dass der Lesegeschwindigkeit große Bedeutung im Worterkennungsprozess zukommt. Wie ausgeprägt und fortgeschritten die Automatisierung in diesem Bereich ist, lässt sich nicht definitiv feststellen. Eine effiziente Worterkennung ist eine Voraussetzung für das Leseverständnis. Erst wenn diese voll ausgebildet ist, stellt das Hörverstehen die begrenzende Leistungsvariable im Verstehensprozess dar. Es ist daher durchaus denkbar, dass

dem Worterkennen – und somit seiner Funktionalität – im Gegensatz zum Hörverstehen diese entscheidende Position im Leseverständnisprozess bei ViertklässlerInnen zukommt, da die Worterkennung noch nicht vollständig automatisiert zu sein scheint. Auch ist zu bedenken, dass die semantische, syntaktische und grammatikalische Komplexität der Sätze des Testverfahrens „Hörverstehen“ aus dem Knuspel-L (Marx, 1998) für Kinder der vierten Schulstufe nur im unteren Leistungsbereich differenziert und Deckeneffekte vorhanden sind (Marx & Jungmann, 2000). Aufgrund des geringen Schwierigkeitsgrades des Testverfahrens (mehr als 73% der SchülerInnen erreichten mindestens 30 der maximalen 34 Punkte) wird ein geringerer Varianzanteil als beim Lesetest erzeugt. Des Weiteren sei darauf hingewiesen, dass die Ergebnisse durchaus mit denen von Tilstra et al. (2009) konform sind. In ihrer Studie ermittelten sie für amerikanische SchülerInnen der vierten Klasse das Dekodieren als wichtigsten Prädiktor für das Leseverständnis, das Hörverstehen brachte weitere 19% Varianzaufklärung für das Leseverständnis. Erst drei Schulstufen höher konnte ein stärkerer Einfluss des Hörverstehens gegenüber dem Worterkennen für SchülerInnen an weiterführenden Schulen nachgewiesen werden. Die vorliegenden Untersuchungsergebnisse bestätigen die Annahmen des „simple view of reading“-Ansatzes dahingehend, dass von einer schriftlich abhängigen Komponente (Worterkennen) und einer schriftlich unabhängigen verbalen Fähigkeitskomponente (Hörverstehen) ausgegangen werden kann.

Für den *Wortschatz* zeigte sich in der untersuchten Stichprobe ein geringer Effekt auf das *Leseverständnis*. Demzufolge ermöglicht ein umfangreicherer Wortschatz ein besseres Leseverständnis. Der Effekt stellte sich jedoch nur im präferierten Modell als statistisch nicht signifikant heraus, was zum einen auf methodische Artefakte zum anderen auf Stichprobeneffekte zurückzuführen ist (siehe Abschnitt 8.3.). In den beiden anderen Modellen bleibt der Einfluss zwar gering, jedoch statistisch signifikant. Angesichts des vorliegenden Ergebnisses kann lediglich darauf geschlossen werden, dass eine gute Leistung im Wortschatztest als Index für allgemeines Wortwissen statistisch gesehen nicht ausschlaggebend für ein gutes Ergebnis im Leseverständnistest sein dürfte. Das verwendete Verfahren zur Erfassung des Wortschatzes scheint nicht geeignet zu sein, die Güte der Qualität der lexikalischen Repräsentationen in seiner Gesamtheit zu erfassen, da hier nur die phonologischen und semantischen Konstituenten des passiven Wortschatzes erhoben wurden und diese eher der verbalen Fähigkeitskomponente (in diesem Fall dem Hörverstehen)

zuzuordnen sind. Desweiteren ist zu argumentieren, dass ein statistisch signifikantes Ergebnis nicht mit einem absolut gültigen Anspruch gleichzusetzen ist, da es lediglich auf Wahrscheinlichkeiten beruht. Im vorliegenden Fall wurde das Signifikanzniveau von 5% knapp verfehlt ($p=.059$), wodurch der Fehler zweiter Art (die fälschliche Annahme, dass der Wortschatz keinen direkten Effekt auf das Leseverständnis hat) relativ hoch ist. Ein statistisch (nicht) signifikantes Ergebnis sagt jedoch nicht zwangsläufig etwas über die praktische Relevanz eines Ergebnisses aus. Es lassen sich empirische Studien und theoretisch begründete Annahmen anführen, die von einem essentiellen Einfluss des Wortschatzes auf das Leseverständnis ausgehen. Tilstra et al. (2009) konnten einen zusätzlichen Anteil an Varianzaufklärung für das Leseverständnis durch verbale Fähigkeiten nachweisen, unabhängig von Dekodieren und Hörverstehen. Bei Kindern der vierten Schulstufe betrug dieser Anteil 5%. Braze et al. (2007) erzielten für ihre Stichprobe ähnliche Ergebnisse. Im Modell von Cromley und Azevedo (2007) besteht für den Wortschatz ein direkter, mittlerer Effekt auf das Verständnis, zusätzlich wird ein geringer Effekt über die Inferenzbildung vermittelt. Wortschatz-Trainingsstudien (z.B. McKeown, Beck, Omanson & Perfetti, 1983) weisen ebenfalls auf die förderliche Wirkung des Wortschatzes bezüglich des Leseverständnisses hin. Auch das National Reading Panel (NICHD, 2000) und das Bundesministerium für Bildung und Forschung (BMBF, 2007) kommen zu dem Schluss, dass ein umfangreicher Wortschatz zu einem guten allgemeinen Leseverständnis beiträgt. McElvany und Schneider (2009) nehmen eine Beeinflussung des Leseprozesses auf allen Ebenen durch den Wortschatz an. Aufgrund der angeführten Argumente scheint es der Autorin der vorliegenden Studie gerechtfertigt, einen eigenständigen, direkten Einfluss des Wortschatzes anzunehmen und somit den „simple view of reading“-Ansatz zur Erklärung des Leseverständnisses um diese Variable zu erweitern.

Die positive Korrelation von mittlerer Stärke zwischen *Wortschatz* und *Hörverstehen* lässt sich dahingehend interpretieren, dass eine reziproke Beziehung zwischen Leistungen im Wortschatztest und Leistungen im Hörverstehenstest bestehen. Mit einem gut ausgeprägten Hörverstehen scheint zum einen ein umfangreicher Wortschatz einher zu gehen, zum anderen kann vermutet werden, dass durch einen umfangreichen Wortschatz mündlich gestellte Instruktionen adäquat aufgenommen und rasch verarbeitet werden können. Dadurch werden vermehrt kognitive Ressourcen für die korrekte Ausführung der Anweisungen frei (vgl.

Aufgabenstellung im Knuspel-L). Folglich sollten mit einem großen Wortschatz bessere Leistungen im Hörverstehen einhergehen, die wiederum einen positiven direkten Effekt auf das Leseverständnis haben. In der durchgeführten Untersuchung wurde der Wortschatz gewissermaßen mündlich erhoben (markieren des entsprechenden Bildes nach lautem Vorlesen der Wörter), sodass sein Einfluss dementsprechend eher auf Seite der Sprachkomponente (Verständnis für gesprochene Sprache) sichtbar ist. Jedoch kann gleichzeitig argumentiert werden, dass der Einfluss des Wortschatzes, unter Annahme der Gültigkeit des *simple view of reading*-Ansatzes, komplett unter die sprachliche Fähigkeitskomponente subsumiert sein sollte. Das vorliegende Ergebnis kann dahingehend interpretiert werden, dass dem Wortschatz respektive dem Wortwissen eine zumindest indirekt vermittelte, eigenständige Funktion im Leseverständnisprozess zukommt. Allerdings muss hier einschränkend erwähnt werden, dass es sich lediglich um korrelative Zusammenhänge zwischen Wortschatz und Hörverstehen handelt, die keine Aussagen über eine kausale Beeinflussung der Variablen zulassen. Aufgrund des theoretischen Rahmens scheint die angeführte Interpretation eine mögliche Erklärung der Ergebnisse zu sein, wobei selbstverständlich der Einfluss weiterer, nicht erhobener, Komponenten durchaus denkbar ist.

In Übereinstimmung mit der „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002) kann argumentiert werden, dass das Leseverständnis als nicht unabhängig von den lexikalischen Fähigkeiten zu betrachten ist. Dem Wortschatz wird neben seinem Beitrag zum Hörverstehen auch eine eigenständige Funktion im Leseverständnisprozess zugeschrieben (vgl. auch Braze et al., 2007). Der starke Einfluss des Worterkennens könnte Perfetti und Hart (2002) zufolge als Beleg für die Kohärenz der Konstituenten gedeutet werden. Sind die lexikalischen Repräsentationen kohärent, so kann eine effiziente (schnelle und bedeutungserfassende) Wortidentifikation stattfinden. Im Wortlesetest (Schabmann et al., 2009) wird die Fähigkeit zum schnellen und fehlerfreien Lesen von Wörtern und Pseudowörtern überprüft. Somit werden Anforderungen an die phonologischen (Graphem-Phonem-Zuordnung), orthographischen (Schreibweise der Wörter) und in gewisser Weise an die semantischen (Erfassen der Bedeutung führt zu weniger Lesefehlern und höherer Geschwindigkeit) Fertigkeiten gestellt. Gute Leistungen im Wortlesetest spiegeln dementsprechend eine hohe lexikalische Qualität der Wortrepräsentationen wider. Infolgedessen könnte interpretiert werden, dass das Wissen über Wörter schon beim Worterkennen wichtig ist und dadurch

ebenfalls auf das Leseverständnis einwirkt. Da vermutet werden kann, dass die Wortschatz-Komponente implizit beim Worterkennen beteiligt ist, ließe sich erklären, dass kein direkter Effekt vom Wortschatz auf das Worterkennen nachgewiesen werden konnte. Eine weitere mögliche Begründung liegt möglicherweise in der bereits erwähnten Messung des Wortschatzes. Die lexikalischen Repräsentationen (Wortschatz) entwickeln sich nach Perfetti und Hart (2002) durch die Erfahrung im Umgang mit der Schriftsprache und der gesprochenen Sprache. Dass die Erfahrung mit der Schriftsprache (repräsentiert durch das Worterkennen) einen Einfluss auf den Wortschatz hat, konnte für die vorliegende Stichprobe nicht bestätigt werden, jedoch besteht ein korrelativer Zusammenhang zwischen Wortschatz und Hörverstehen. Es scheint sich das Hörverstehen/der Wortschatz durch den Umgang mit der Sprache zu entwickeln, was gleichzeitig positiv mit dem sprachlichen Wissen/der Fähigkeiten im Hörverstehen zusammenhängt.

Es lässt sich zusammenfassend festhalten, dass die Varianz im *Leseverständnis* zu 61% durch die direkten Effekte von *Worterkennen*, *Wortschatz* und *Hörverstehen*, sowie der Korrelation zwischen *Wortschatz* und *Hörverstehen* erklärt werden kann. Die Ergebnisse bestätigen einerseits die Annahmen des „simple view of reading“-Ansatzes (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000), wobei dem *Worterkennen* eine größere Bedeutung für das *Leseverständnis* zukommt als dem *Hörverstehen*. Andererseits ist der *Wortschatz* positiv mit der schriftlich unabhängigen verbalen Fähigkeitskomponente verbunden und dem Wortschatz wird zusätzlich eine direkte Funktion im Leseverständnisprozess zugesprochen. Darüber hinaus wäre es denkbar, dass lexikalische Repräsentationen hoher Qualität mit dem Worterkennen und dem Leseverständnis positiv zusammenhängen. Die Varianz im Leseverständnis scheint demnach nicht ausschließlich durch die Komponenten des „simple view of reading“-Ansatzes erklärbar, sodass eine Erweiterung um die Variable Wortschatz in Anlehnung an die „lexical quality hypothesis“ (Perfetti & Hart, 2001, 2002) sinnvoll erscheint.

8.2. Entwicklungsaspekte des Leseverständnisses bei Grundschulkindern

Dieses Kapitel wurde gemeinsam von Stefanie Dorn (2011), Yvonne Huemer (2011) und Maria Klausecker (2011) verfasst.

Im Rahmen des Gesamtprojektes zur Untersuchung des Leseverständnisses in Abhängigkeit von den Prädiktoren *Worterkennen*, *Hörverstehen* und *Wortschatz* sind drei zeitgleiche Parallelstudien entstanden, die sich lediglich in Hinblick auf die untersuchten Klassenstufen unterscheiden. Um eine Vergleichbarkeit der Ergebnisse der zweiten (Klausecker, 2011), dritten (Huemer, 2011) und vierten Schulstufe (siehe Kapitel 7) zu gewährleisten, wurden im Rahmen dieses Projektes gleiche Zielsetzung und Fragestellungen (siehe Kapitel 5), Testverfahren (siehe Kapitel 6.2.) und Auswertungsmethoden (siehe Kapitel 6.3.) verwendet.

Die Analyse mittels Strukturgleichungsmodellen betont die wichtige Rolle des *Worterkennens*, *Hörverstehens* und *Wortschatzes* für das *Leseverständnis* über alle untersuchten Klassenstufen hinweg. In der zweiten Klasse wird das *Leseverständnis* am stärksten durch das *Worterkennen* beeinflusst. Statistisch signifikante Effekte werden auch für den *Wortschatz* und das *Hörverstehen* postuliert, wobei dem *Wortschatz* eine entscheidendere Funktion als dem *Hörverstehen* im Leseverständnisprozess zuzukommen scheint. Auch in der dritten Klasse erweist sich das *Worterkennen* als wichtigster der drei Prädiktoren, gefolgt von *Hörverstehen* und *Wortschatz*, deren Einfluss in etwa gleich stark ausgeprägt ist. Am Ende der Grundschulzeit wird das *Leseverständnis* essentiell durch das *Worterkennen* bedingt. Als zweit wichtigster Prädiktor kann das *Hörverstehen* angesehen werden, wohingegen der direkte Effekt des *Wortschatzes* von eher untergeordneter Funktion sein dürfte.

Summa summarum kann dem *Worterkennen* die tragende Rolle im Leseverstehensprozess bei Kindern der zweiten, dritten und vierten Schulstufe zugesprochen werden. Dieses Ergebnis steht unter anderem im Einklang mit den Befunden von Braze et al. (2007) und Tilstra et al. (2009). Auch Marx und Jungmann (2000) betonen die Relevanz des Worterkennens für das Leseverständnis über die vierte Klasse hinaus. Für die Fähigkeiten im Worterkennungsprozess scheint insbesondere die Lesegeschwindigkeit verantwortlich zu sein, was als Hinweis auf zunehmende Automatisierung bei der Wortidentifikation angesehen werden kann (vgl. auch Klicpera & Gasteiger-Klicpera, 1998; Landerl & Willburger, 2009; siehe auch Interpretation

Dorn, 2011). Aus der Zunahme der Lesegeschwindigkeit resultiert ein verbessertes Leseverständnis (u.a. Hosp & Fuchs, 2005), was sich auch schon bei SchülerInnen der zweiten Schulstufe beobachten lässt (siehe Klausecker, 2011). Dies spiegelt auch die Unterrichtsmethode „synthetic phonics“ (vgl. Klicerpa et al., 2010) wider, welche die Konsistenz der Graphem-Phonem-Zuordnung der deutschen Sprache ausnutzt, sodass schon bei LeseanfängerInnen, zu denen Kinder der zweiten Klasse zuzuordnen sind, von einer sicheren und gleichzeitig raschen Wortidentifikation ausgegangen werden kann (Klicpera & Gasteiger-Klicpera, 1993). Die Befunde von Huemer (2011) und die Darstellung in Kapitel 7. und Kapitel 8.1. legen die Schlussfolgerung nahe, dass eine effiziente Worterkennung die Basis für alle weiteren Leseverständnisprozesse darstellt.

Der Einfluss des Hörverstehens steigt über alle untersuchten Klassenstufen hinweg kontinuierlich an. In der dritten und vierten Klasse stellt das Hörverstehen den zweit wichtigsten Prädiktor des Leseverständnisses dar. In der zweiten Klasse ist der direkte Effekt des Hörverstehens als geringer zu betrachten, jedoch lässt sich ein statistisch signifikanter Beitrag des Hörverstehens zum Leseverständnis nachweisen. Dieser Befund geht Hand in Hand mit den Ergebnissen von Rost und Hartmann (1992). Auch belegt die Analyse der vorliegenden Daten eine generelle Überlegenheit des Worterkennens im Gegensatz zum Hörverstehen, was sich mit den Ergebnissen von Tilstra et al. (2009) deckt.

In der zweiten Klasse kommt dem Wortschatz nach dem Worterkennen eine wesentliche Bedeutung im Leseverständnisprozess zu. Ein größerer Wortschatz trägt maßgeblich zum erfolgreichen Textverständnis bei und der Leseprozess wird zusätzlich auf allen Ebenen vom Wortschatz beeinflusst (u.a. BMBF, 2007; McElvany & Schneider, 2009; NICHD, 2000). Im Rahmen dieses Projektes zeigt sich jedoch eine stetige Abnahme der Beeinflussung des Leseverständnisses durch den Wortschatz in den höheren Klassen. Für einen möglichen Erklärungsversuch siehe Kapitel 8.1. und 8.3.

Zusammenfassend können sich die Ergebnisse als Beleg für den „simple view of reading“-Ansatz deuten lassen. In allen drei Klassenstufen konnte der Beitrag der zwei Komponenten Worterkennen und Hörverstehen für das Leseverständnis nachgewiesen werden. Dennoch stellt der Wortschatz eine wichtige Einflussgröße dar, sodass eine Berücksichtigung dieser

Variable vorteilhaft erscheint, um der Komplexität des Leseverständnisprozesses ansatzweise gerecht zu werden.

8.3. Kritik

Verfasst von Stefanie Dorn (2011).

Kritische Anmerkungen an der durchgeführten Studie im Rahmen dieser Diplomarbeit betreffen zum einen die Stichprobe, zum anderen die verwendeten Testverfahren. Die vorliegende Stichprobe besteht aus 113 Kindern der vierten Volksschulklasse in ländlichen Gebieten in Oberösterreich. Es ist hierbei zu bedenken, dass im Vergleich zu Schulen in städtischen Gebieten der Anteil an ausländischen Kindern (andere Muttersprache als Deutsch) geringer ist und Unterschiede in den sozioökonomischen Rahmenbedingungen in Stadt und Land bestehen (vgl. Statistik Austria, 2011). Auch fand eine „Selbstselektion“ der UntersuchungsteilnehmerInnen statt. Die Eltern mussten die freiwillige Teilnahme ihrer Kinder bewilligen. Es ist davon auszugehen, dass insbesondere leistungsstarke Kinder beziehungsweise deren Eltern, Interesse an einer Teilnahme an dieser Untersuchung zeigten. Es ist daher nicht auszuschließen, dass die Ergebnisse die Leistungsfähigkeit der Population überschätzen. Somit ist eine Allgemeingültigkeit der Befunde nicht nur im Hinblick auf sozioökonomische Unterschiede in Stadt und Land einzuschränken, sondern auch auf das mögliche Problem der Selbstselektion der Stichprobe (tendenziell eher leistungsstarke Kinder im Bereich Lesen in der vierten Volksschulklasse).

Bezüglich der notwendigen Stichprobengröße zur Berechnung von Strukturgleichungsmodellen bestehen unterschiedliche Meinungen in der Literatur. Es wird zumeist ein Minimalumfang von 100 bis 150 empfohlen (Buch, 2007). Das Stichprobenminimum von 100 konnte in dieser Arbeit realisiert werden, wobei eine größere Stichprobe für effizientere (aussagekräftigere) Ergebnisse jedoch wünschenswert gewesen wäre (Marsh, Hau, Balla & Grayson, 1998).

Für die Erhebung der latenten Variablen *Worterkennen*, *Wortschatz*, *Hörverstehen* und *Leseverständnis* wurde je ein Indikator eingesetzt. Es kann zwar davon ausgegangen werden, dass die standardisierten Testverfahren die latenten Variablen hinreichend gut messen, allerdings wären zumindest drei Indikatoren pro Variable für eine reliable und adäquate

Erfassung erstrebenswert (Marsh et al., 1998). Dies konnte aus Mangel an Verfügbarkeit geeigneter Testverfahren, insbesondere für den deutschsprachigen Raum und im Hinblick auf die große Altersspanne bei der Gesamterhebung (2.-4. Schulstufe, siehe Huemer 2011; Klausecker, 2011), sowie aus testtheoretischen Gründen der Ökonomie und Zumutbarkeit, nicht realisiert werden.

Das Leseverständnis wurde durch den Subtest „Textverständnis“ des ELFE 1-6 (Lenhard & Schneider, 2006) erhoben. Da das Leseverständnis nicht nur die Textebene betrifft, sollte für eine umfassende Erhebung ergänzend auch die Wort- und Satzebene berücksichtigt werden. Dies könnte durch die komplette Vorgabe des ELFE 1-6 realisiert werden, wodurch gleichzeitig drei Indikatoren für die Erfassung der latenten Variable vorhanden wären. Eventuell würde sich so auch ein verstärkter Einfluss des Wortschatzes auf das Leseverständnis nachweisen lassen (vgl. McElvany & Schneider, 2009).

Der Knuspel-L (Marx, 1998) stellt ein Verfahren dar, das bereits ab Ende der ersten Klasse einsetzbar ist. Wie bereits im Abschnitt 8.1. (Interpretation) angesprochen, scheint der Subtest „Hörverstehen“ für Kinder der vierten Klasse im oberen Leistungsbereich nicht mehr ausreichend zu differenzieren, sodass durch verfahrensbedingte Deckeneffekte eine Varianzeinschränkung bedingt sein könnte. Auch betonen Hoover und Gough (1990), dass inhaltlich parallele Testverfahren (korrespondierender Inhalt, z.B. Einsatz von Erzähltexten) zur Erfassung des Hör- und Leseverständnisses notwendig sind, um den „simple view of reading“-Ansatz zu überprüfen. Dies könnte in zukünftigen Studien durch die zusätzliche Vorgabe des Subtests „Leseverstehen“ umgesetzt werden.

Um das Ausmaß des bestehenden Wortwissens ganzheitlich zu erfassen, wären zumindest zwei komplementäre Wortschatzmessungen erforderlich (Protopapas et al., 2007). Der hier verwendete PPVT-III (Dunn & Dunn, 1997) bildet lediglich den passiven Wortschatz ab und kann als Indikator für mündliche Sprachfertigkeiten angesehen werden. Die Ergänzung durch aktive Wortschatztests, die beispielsweise das Verständnis für Wortbedeutungen überprüfen, eine Definition von vorgegebenen Wörtern oder das Benennen von abgebildeten Gegenständen verlangen (vgl. Joshi, 2005), wäre zukünftig anzustreben. Da sich die lexikalischen Repräsentationen unter anderem durch die Erfahrung mit der Schriftsprache entwickeln, könnte der Wortschatz im Sinne der „lexical quality hypothesis“ (Perfetti & Hart,

2001, 2002) durch Verfahren erhoben werden, welche die Leseerfahrung überprüfen. Braze et.al. (2007) haben dies beispielsweise realisiert indem den Testpersonen Listen mit fiktiven Autoren und Magazinen zur Beurteilung bezüglich deren realer Existenz vorgelegt wurden.

Des Weiteren ist kritisch anzumerken, dass durch die verwendete adaptierte Version des PPVT-III für Kinder der vierten Klassen ebenfalls Deckeneffekte beobachtbar sind und somit der erzeugte Varianzanteil des Wortschatzes eher als gering zu beurteilen ist. Um die Vergleichbarkeit dieser Stichprobe mit der jüngerer Kinder zu gewährleisten (vgl. Huemer, 2011; Klausecker, 2011), wurden zu Beginn der Wortschatztestung 12 Items vorgelegt, die für Kinder im Alter von sechs bis sieben Jahren vorgesehen, und somit als zu leicht zu bewerten sind. Da es sich bei der Stichprobe um tendenziell eher gute Leser handeln dürfte („Selbstselektion“), sind auch die 24 Items, die für einen Altersbereich zwischen 12 und 16 Jahren ausgelegt sind, nicht in der Lage ausreichend zwischen den Leistungen zu differenzieren. Für eine genauere Messung des passiven Wortschatzes bei ViertklässlerInnen wäre es daher notwendig gewesen, weitere Sets in aufsteigender Schwierigkeit vorzugeben. Vor allem für Kinder der zweiten Schulstufe wäre jedoch die damit verbundene Verlängerung der Testdauer und die steigenden Anforderungen an die Konzentration und Leistungsfähigkeit als grenzwertig zu betrachten (siehe Klausecker, 2011), sodass von der Vorgabe weiterer Items abgesehen wurde.

Ein letzter Kritikpunkt liegt im nicht berücksichtigten Einfluss des Leseverständnisses auf die latenten Variablen beziehungsweise in der nicht berücksichtigten Endogenität derselben in Hinblick auf das Leseverständnis. Für eine ausführliche Diskussion zum Thema Endogenität erklärender Variablen und deren Konsequenzen für die Konsistenz der Schätzergebnisse sei z.B. auf Greene (2008) verwiesen.

8.4. Ausblick und Relevanz

Verfasst von Stefanie Dorn (2011).

In dieser Untersuchung wurden das *Worterkennen*, der *Wortschatz* und das *Hörverstehen* als determinierende Variablen für das *Leseverständnis* identifiziert. Angesichts der angeführten Kritikpunkte wäre für zukünftige Replikationen eine umfangreichere Erfassung dieser Variablen bei einer größeren Stichprobe wünschenswert. Auch die Überprüfung eines gemeinsamen Modells des Leseverständnisses für SchülerInnen der 2. bis 4. Grundschulstufe ist naheliegend. Des Weiteren wäre für die vorliegende Fragestellung ein Vergleich zwischen städtischen und ländlichen Regionen unter Berücksichtigung soziostruktureller (z.B. Migrationsstatus, sozioökonomischer Status) und prozessbezogener Merkmale (z.B. Bildungsressourcen, Besitz von Kulturgütern) interessant, da für diese positive Zusammenhänge mit der Schulleistung und der Kompetenzentwicklung des Lesens bestehen (für einen Überblick siehe Schaffner, 2009).

Die Berücksichtigung weiterer Bedingungsfaktoren des Leseverständnisses wie etwa kognitive Faktoren (Intelligenz, verbales Arbeitsgedächtnis, Vorwissen oder metakognitives Strategiewissen), motivationale Faktoren (habituelle und aktuelle Lesemotivation oder das thematische Interesse), Lesehäufigkeit (siehe Schaffner, 2009) und Geschlecht (Schwantner & Schreiner, 2010) könnte ebenfalls ein Ziel künftiger Studien sein. Auch wäre zu überlegen, ob der Einfluss des Leseverständnisses auf die Prädiktoren (insbesondere Wortschatz) überprüft wird, da zum Beispiel McElvany und Schneider (2009) oder auch das Bundesministerium für Bildung und Forschung (2007) annehmen, dass zwischen Leseverständnis- und Wortschatzentwicklung eine reziproke Beziehung besteht.

Ein weiterer Aspekt, der im 21. Jahrhundert an enormer Bedeutung gewonnen hat, ist im sinnerfassenden Lesen elektronischer Medien (Internet, eMails, Blogs) zu sehen. Der größte Unterschied zu gedruckten Medien besteht in der dynamischen Darstellungsform dieser Texte und auch die Anforderungen an die Fähigkeiten des Lesers/der Leserin haben sich gewandelt. So muss auf Texte zugegriffen, diese verstanden und aufeinander bezogen, und letztendlich auch die Glaubwürdigkeit und Gültigkeit bewertet werden können. Die Ergebnisse der PISA-Studie 2009 (OECD, 2010) geben Anlass zu der Vermutung, dass das Leseverständnis für

elektronische Texte schlechter als für gedruckte Texte ist (Schwantner & Schreiner, 2010). Im Rahmen zukünftiger Studien ist es daher wichtig, die wesentlichen Komponenten des sinnerfassenden Lesens elektronischer Medien adäquat zu erfassen und im Hinblick auf mögliche Förderungsmaßnahmen zu evaluieren.

Ein umfassendes Modell des Leseverständnisses, das sämtliche Einflussfaktoren berücksichtigt, dürfte aufgrund der Komplexität des Lese(verständnis)prozesses und der Individualität jedes Lesers/jeder Leserin nur schwer zu erstellen sein. Dies kann auch nicht als vorrangiges Ziel dieser Studie betrachtet werden. Entscheidend war es viel eher die wesentlichen Komponenten des Leseverständnisses zu eruieren und somit geeignete Ansatzpunkte zur Förderung der Lesekompetenz in der Praxis aufzuzeigen. Dass Förderung nicht nur für sogenannte RisikoschülerInnen notwendig ist, sondern ein allgemeines Anliegen des schulischen Unterrichts sein sollte, hat die PISA-Studie 2009 eindrucksvoll bewiesen. Allerdings ist hier zu betonen, dass auch der Familie eine entscheidende Funktion bei der Förderung zukommt (McElvany & Schneider, 2009).

Als essentieller Faktor des Leseverständnisses hat sich das Worterkennen herausgestellt. Demnach kommt der Förderung der basalen Lesefertigkeiten eine ausschlaggebende Bedeutung zu, da diese die Voraussetzungen für höhere Lesefertigkeiten darstellen und somit letztendlich auch die Voraussetzungen für das Lernen bilden (McElvany & Schneider, 2009). Als Ansatzpunkt zur Förderung bieten sich insbesondere Trainings der Leseflüssigkeit (z.B. „guided repeated oral reading“, McElvany & Schneider, 2009, „paired repeated reading, NICHD, 2000) an, die sich im Unterrichtskontext leicht realisieren lassen. Der ELFE 1-6 (Lenhard & Schneider, 2006) liegt zusätzlich zur Papierversion in computerisierter Form vor. Die damit verbundenen Vorteile (z.B. automatische Auswertung, Standardisierung, Interpretationsvorschläge, Ökonomie, etc.) werden durch die adaptive Itemvorgabe ergänzt. Auch ist durch die PC-Vorgabe eine Schwellenmessung der visuellen Worterkennung (Lesegeschwindigkeitstest) durchführbar. Als weiterer positiver Aspekt der PC-Fassung ist die Möglichkeit zur gezielten Förderung des Leseverständnisses zu sehen. Die Ergebnisse aus dem Lesetest können in das Trainingsprogramm eingelesen werden und der Computer wählt dem Leistungsprofil des Kindes entsprechende Übungen aus (Lenhard & Lenhard, 2009a).

In dieser Untersuchung wurde explizit der Einfluss des Wortschatzes für das Leseverständnis betont. Unabhängig von den ermittelten statistischen Effekten ist hier ein weiterer wesentlicher Ansatz zur Förderung des Leseverständnisses zu sehen. Exemplarisch sei auf „Drop everything and write“-Programm (Joshi, Aaron, Hill, Ocker, Boulware-Gooden & Rupley, 2008) hingewiesen. Dieser Ansatz sieht vor, dass Kinder täglich 20 Minuten lang nur schriftlich miteinander kommunizieren. Die positiven Auswirkungen konnten empirisch bestätigt werden. Für weitere Förderprogramme und Methoden, die auch auf die Erhöhung der allgemeinen Lesekompetenz abzielen, sei auf den Bericht des National Reading Panels (NICHD, 2000), die Expertise des BMBF (2007), Klicpera et al. (2010) und Lenhard und Schneider (2009) verwiesen.

9. Zusammenfassung

Verfasst von Stefanie Dorn (2011).

9.1. Deutsche Zusammenfassung

Das Leseverständnis wird dem „simple view of reading“-Ansatz (Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000) zufolge durch die zwei Komponenten Dekodieren (Worterkennen) und Hörverstehen determiniert. Darüber hinaus werden lexikalischen Fertigkeiten eine entscheidende Funktion bei Prozessen des Leseverständnisses zugesprochen (vgl. „lexical quality hypothesis“, Perfetti & Hart, 2001, 2002; Braze et al., 2007; Tilstra et al., 2009; Protopapas et al., 2007). Das Ziel der vorliegenden Studie besteht in der Identifikation des Beitrags der Komponenten Worterkennen, Wortschatz und Hörverstehen für das Leseverständnis bei Kindern der vierten Schulstufe. Die zeitgleich entstandenen Parallelstudien von Klausecker (2011) und Huemer (2011) thematisieren diesen Sachverhalt bei Kindern der zweiten, beziehungsweise dritten Schulstufe.

113 SchülerInnen der vierten Klasse in Oberösterreich wurden vier standardisierte Testverfahren zur Erfassung der einzelnen Komponenten vorgelegt. Um die latenten Variablen Worterkennen, Wortschatz, Hörverstehen und Leseverständnis zu operationalisieren, wurden ein Wortschatztest (PPVT-III; Dunn & Dunn, 1997), ein Hörverstehenstest (Knuspel-L; Marx, 1998), ein Textverständnistest (ELFE 1-6; Lenhard & Schneider, 2006) und ein Verfahren, das die Fähigkeiten im Worterkennen überprüft (Wortlesetest; Schabmann et al., 2009), eingesetzt.

Die Datenanalyse mittels Strukturgleichungsmodellen ergab eine Varianzaufklärung des Leseverständnisses von insgesamt 61% durch die drei Variablen. Das Ergebnis deutet darauf hin, dass den beiden Komponenten des „simple view of reading“-Ansatzes die wichtigste Funktion für das Leseverständnis zukommt, wobei dem Worterkennungsprozess der größte Beitrag zugeschrieben wird. Der direkte Effekt des Wortschatzes ist gering, wobei gleichzeitig eine moderate Korrelation mit dem Hörverstehen besteht, sodass vermutet werden kann, dass lexikalische Fertigkeiten eher indirekt, aber dennoch in einem positiven Zusammenhang mit Verständnisprozessen stehen. Obwohl sich der direkte Einfluss des

Wortschatzes auf das Leseverständnis in der untersuchten Stichprobe als statistisch nicht signifikant erweist, scheinen lexikalische Repräsentationen hoher Qualität im Worterkennungsprozess bedeutsam zu sein. Auch wird lexikalischen Fertigkeiten allgemein eine ausschlaggebende Rolle im Leseprozess zugeschrieben (vgl. BMBF, 2007; Cromley & Azevedo, 2007; McElvany & Schneider, 2009; NICHD, 2000). Daher wird argumentiert, dass der Wortschatz als ein wesentlicher Faktor in einem Modell des Leseverständnisses berücksichtigt werden sollte. Desweiteren ist anzuführen, dass unabhängig von den vorliegenden statistischen Ergebnissen eine praktische Relevanz für alle drei Variablen besteht, da sie geeignete Ansatzpunkte zur Förderung des Leseverständnisses darstellen.

Für weitere Studien empfiehlt sich eine umfangreichere Erfassung der Variablen an einer größeren und gleichzeitig heterogeneren Stichprobe. Auch erscheint eine Berücksichtigung weiterer Einflussfaktoren wie etwa kognitive, motivationale oder sozioökonomische Determinanten, für ein umfassenderes Modell des Leseverständnisses sinnvoll.

Schlüsselwörter: „simple view of reading“-Ansatz, “lexical quality hypothesis”, Leseverständnis, Worterkennen, Hörverstehen, lexikalische Fähigkeiten, Wortschatz

9.2. Abstract

The existing literature provides some evidence that reading comprehension simply consists of two components, namely decoding and linguistic comprehension (“simple view of reading”; Gough & Tunmer, 1986; Hoover & Gough, 1990; Marx & Jungmann, 2000). In addition, a lot of research emphasizes the unique, important role of lexical skills that cause to a considerable extent the ability in reading comprehension over and above the “simple view of reading” (“lexical quality hypothesis”, Perfetti & Hart, 2001, 2002; Braze et al., 2007; Tilstra et al., 2009; Protopapas et al., 2007; among others). The purpose of the present study is to identify the contribution of each of the components mentioned above on reading comprehension based on a sample of elementary school pupils in fourth grade. Simultaneously to this thesis Klausecker (2011) and Huemer (2011) address the same issues in parallel studies for pupils in the second and third grade, respectively.

113 fourth-graders in Upper Austria were tested on four standardized tasks used as instruments for the latent variables decoding, vocabulary, linguistic comprehension and reading comprehension: Word and pseudoword reading accuracy and fluency (Wortlesetest, Schabmann et al., 2009), receptive vocabulary (PPVT-III, Dunn & Dunn, 1997), listening comprehension (Knuspel-L; Marx, 1998) and reading comprehension (ELFE 1-6, Lenhard & Schneider, 2006).

Using structural equation models for the data analysis, a total of 61% of the variance in reading comprehension can be explained by these measurements of the three latent variables. The two “simple view of reading”-components seem to have more predictive strength in explaining differences in reading comprehension, with word reading skills having the highest impact. The direct effect of vocabulary on reading comprehension is low but the moderate correlation found between vocabulary and listening comprehension can be interpreted as a merely indirect, but positive relation between vocabulary and comprehension processes. Even though the results do not detect a statistically significant impact of vocabulary on reading comprehension for the sample used in this study, one can argue that high quality lexical representations are fundamental for skilled performance in decoding. Additionally, lexical skills are often regarded as a relevant factor for reading abilities (c.f. BMBF, 2007; Cromley & Azevedo, 2007; McElvany & Schneider, 2009; NICHD, 2000) and should therefore be included to a model of reading comprehension. Regardless of the statistically unique contributions of each component, some practical implication of the findings should be highlighted: Word reading skills, linguistic comprehension and lexical skills each can be regarded as a crucial basis for programs enhancing reading comprehension.

For future research it seems necessary to assess the key components more comprehensively and in a sample consisting of more heterogeneous pupil. Additionally, other related variables like cognitive, motivational and socioeconomic factors should be taken into account for modeling reading comprehension.

Keywords: “simple view of reading”, “lexical quality hypothesis”, decoding, lexical skills, vocabulary, linguistic comprehension, listening comprehension, reading comprehension

10. Literaturverzeichnis

- Abbott, V., Black, J. B. & Smith, E. E. (1985). The representation of scripts in memory. *Journal of Memory and Language* 24 (2), 179–199.
- Backhaus, K., Erichson, B., Plinke, W. & Weiber, R. (2008). *Multivariate Analysemethoden. Eine anwendungsorientierte Einführung* (12. vollständig überarbeitete Auflage). Springer: Berlin, Heidelberg, New York.
- Bartlett, F. C. (1932). *Remembering: A study in experimental and social psychology*. Oxford: University press.
- Beck, I. L., Perfetti, C. A. & McKeown, M. G. (1982). Effects of long-term vocabulary instruction on lexical access and reading comprehension. *Journal of Educational Psychology*, 74 (4), 506–521.
- Bell, L. C. & Perfetti, C. A. (1994). Reading skill: Some adult comparison. *Journal of Educational Psychology*, 86 (2), 244–255.
- Besner, D. & Roberts, M. A. (2003). Reading nonwords aloud: Results requiring change in the dual route cascaded model. *Psychonomic Bulletin & Review*, 10 (2), 398–404.
- Bourassa, D. C., Levy, B. A., Dowin, S. & Casey, A. (1998). Transfer effects across contextual and linguistic boundaries: Evidence from poor readers. *Journal of Experimental Child Psychology*, 71 (1), 54–61.
- Braze, D., Tabor, W., Shankweiler, D. P. & Mencl, E. W. (2007). Speaking up for vocabulary: Reading skill differences in young adults. *Journal of Learning Disabilities*, 40 (3), 226–243.
- Buch, S. (2007). Strukturgleichungsmodelle-Ein einführender Überblick (ESCP-EAP Working Paper No. 29). Berlin: Europäische Wirtschaftshochschule. Zugriff am 18.11.2010. Verfügbar unter http://www.escp-eap.eu/uploads/media/Buch_Strukturgleichungsmodelle_ESCP_EAP_WP29_04.pdf

- Bundesministerium für Bildung und Forschung (BMBF) (Hrsg.). (2007). *Förderung der Lesekompetenz-Expertise* (Bildungsforschung Band 17). Bonn/ Berlin: Bundesministerium für Bildung und Forschung. Zugriff am 20.01.2011. Verfügbar unter <http://www.bmbf.de>
- Byrne, B. M. (2001). *Structural equation modeling with AMOS: Basic concepts, applications, and programming*. New Jersey, NJ: Erlbaum.
- Cain, K., Oakhill, J. & Lemmon, K. (2004). Individual differences in the inference of word meanings from context: The influence of reading comprehension, vocabulary knowledge and working memory capacity. *Journal of Educational Psychology*, 96 (4), 671–681.
- Chen, R. S. & Vellutino, F. R. (1997). Prediction of reading ability: A cross-validation study of the simple view of reading. *Journal of Literacy Research*, 29 (1), 1–24.
- Christmann, U. & Groeben, N. (1999). Psychologie des Lesens. In B. Franzmann, K. Hasemann, D. Löffler & E. Schön (Hrsg.), *Handbuch Lesen* (S. 145–223). München: K.G. Sauer.
- Coltheart, M. (1978). Lexical access in simple reading tasks. In G. Underwood (Ed.), *Strategies of information processing* (pp. 112–174). New York: Academic Press.
- Coltheart, M. & Rastle, K. (1994). Serial processing in reading aloud: Evidence for dual-route models of reading. *Journal of Experimental Psychology: Human Perception and performance*, 20 (6), 1197–1211.
- 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), 204–256.
- Cromley, J. G. & Azevedo, R. (2007). Testing and refining the direct and inferential mediation model of reading comprehension. *Journal of Educational Psychology*, 99 (2), 311–325.
- Curtis, M. (1980). Development of components of reading skill. *Journal of Educational Psychology*, 72 (5), 656–669.

- de Jong, P. F. & van der Leij, A. (2002). Effects of phonological abilities and linguistic comprehension on the development of reading. *Scientific Studies of Reading*, 6 (1), 51–77.
- Dreyer, L. G. & Katz, L. (1992). An examination of “the simple view of reading”. In C. K. Kinzer & D. J. Leu (Eds.), *Literacy research, theory, and practice: Views from many perspectives. Forty-first yearbook of the National Reading Conference* (pp. 169–175). Washington, DC: The National Reading Conference, Inc.
- Dunn, L. M. & Dunn, L. M. (1997). *Peabody Picture Vocabulary Test (PPVT-III)*. 3rd edition. Circle Pines, MN: American Guidance Service.
- Ehri, L. C. (1997). Sight word learning in normal readers and dyslexics. In B. Blachman (Ed.), *Foundations and reading acquisition and dyslexia* (pp. 163–190). London: Erlbaum.
- Elbro, C. (1996). Early linguistic abilities and reading development: A review and a hypothesis. *Reading and Writing: An Interdisciplinary Journal*, 8, 453–485.
- Frith, U. (1985). Beneath the surface of developmental dyslexia. Are comparisons between developmental and acquired disorders meaningful? In K. E. Patterson, J. C. Marshall & M. Coltheart (Eds.), *Surface dyslexia: Neuropsychological and cognitive studies of phonological reading* (pp. 301–330). London: Erlbaum.
- Fuchs, L. S., Fuchs, D., Hosp, M. K. & Jenkins, J. R. (2001). Oral reading fluency as an indicator of reading comprehension: A theoretical, empirical, and historical analysis. *Scientific Studies of Reading*, 5 (3), 239–256.
- Gough, P. B. & Tunmer, W. E. (1986). Decoding, reading, and reading disability. *Remedial and Special Education*, 7 (1), 6–10.
- Graesser, A. C., Singer, M. & Trabasso, T. (1994). Constructing inferences during narrative text comprehension. *Psychological Review*, 101 (3), 371–395.
- Graf, E. (1994). *Lese-Rechtschreib-Schwäche. Ein prozessanalytischer Ansatz*. Bern: Peter Lang.
- Greene, W. H. (2008). *Econometric analysis* (6. ed.). Upper Saddle River, NJ: Pearson Prentice Hall.

- Grisseemann, H. (1996). *Von der Legasthenie zum gestörten Schriftspracherwerb*. Bern: Huber.
- Günther, K. B. (1986). Ein Stufenmodell der Entwicklung kindlicher Lese- und Schreibstrategien. In H. Brügelmann (Hrsg.), *ABC und Schriftsprache: Rätsel für Kinder, Lehrer und Forscher* (S. 32–54). Konstanz: Faude.
- Guindon, R. & Kintsch, W. (1984). Priming macropropositions: Evidence for the primacy of macropropositions in the memory for text. *Journal of Verbal Learning and Verbal Behavior*, 23 (4), 508–518.
- Hagtvet, B. E. (2003). Listening comprehension and reading comprehension in poor decoders: Evidence for the importance of syntactic and semantic skills as well as phonological skills. *Reading and Writing: An Interdisciplinary Journal*, 16, 505–539.
- Heidemann-Menda, W. (2006). *Zum Leseverständnis lernbehinderter Sonderschüler. Theorie – Ist-Zustand – Fördermöglichkeiten*. Sankt Augustin: Autor.
- Hoover, W. A & Gough, P. B. (1990). The simple view of reading. *Reading and Writing: An Interdisciplinary Journal*, 2, 127–160.
- Hosp, M. & Fuchs, L. (2005). Using CBM as an indicator of decoding, word reading, and comprehension: Do the relations change with grade? *School Psychology Review*, 34 (1), 9–26.
- Hoyle, R. H. (1995). Structural equation modeling approach – Basic concepts and fundamental issues. In R. H. Hoyle (Ed.), *Structural equation modeling – Concepts, issues and applications* (pp. 1–15). Thousand Oaks, CA: Sage.
- Huemer, Y. (2011). *Prädiktoren des Leseverständnisses bei Grundschulkindern der dritten Schulstufe – Erweiterung des „simple view of reading“ Ansatzes?* Unveröffentlichte Diplomarbeit, Universität, Wien.
- Joshi, R. M. (2005). Vocabulary: A critical component of comprehension. *Reading & Writing Quarterly*, 21 (3), 209–219.
- Joshi, R. M. & Aaron, P. G. (2000). The component model of reading: Simple view of reading made a little more complex. *Reading Psychology*, 21 (2), 85–97.

- Joshi, R. M., Aaron, P.G., Hill, N., Ocker E., Boulware-Gooden, R. & Rupley, W. (2008). Drop everything and write (DEAW): An innovative program to improve literacy skills. *Learning Inquiry*, 2 (1), 1–12.
- Kintsch, W. (1974). *The representation of meaning in memory*. Hillsdale/ NJ: Laurence Erlbaum.
- Kintsch, W. & van Dijk, T. A. (1978). Toward a model of text comprehension and production. *Psychological Review*, 85 (5), 363–394.
- Kirschhock, E.-M. (2004). *Entwicklung schriftsprachlicher Kompetenzen im Anfangsunterricht*. Bad Heilbrunn: Julius Klinkhardt.
- Klausecker, M. (2011). *Prädiktoren des Leseverständnisses bei Grundschulkindern der zweiten Schulstufe – Erweiterung des „simple view of reading“ Ansatzes?* Unveröffentlichte Diplomarbeit, Universität Wien
- Klicpera, C. & Gasteiger-Klicpera, B. (1993). *Lesen und Schreiben- Entwicklung und Schwierigkeiten. Die Wiener Längsschnittuntersuchung über die Entwicklung, den Verlauf und die Ursachen von Lese- und Schreibschwierigkeiten in der Pflichtschulzeit*. Bern: Huber.
- Klicpera, C. & Gasteiger-Klicpera, B. (1995). *Psychologie der Lese- und Schreibschwierigkeiten. Entwicklung, Ursachen, Förderung*. Weinheim: Beltz.
- Klicpera, C. & Gasteiger-Klicpera, B. (1998). *Psychologie der Lese- und Schreibschwierigkeiten. Entwicklung, Ursachen, Förderung* (2. Aufl.). Weinheim: Beltz.
- Klicpera, C., Schabmann, A. & Gasteiger-Klicpera, B. (2003). *Legasthenie. Modelle, Diagnose, Therapie und Förderung*. München: Reinhardt.
- Klicpera, C., Schabmann, A. & Gasteiger-Klicpera, B. (2010). *Legasthenie – LRS. Modelle, Diagnose, Therapie und Förderung* (3., aktualisierte Auflage). München: Ernst Reinhardt.
- Kline, R.B. (1998). *Principles and practice of structural equation modeling*. New York, NY: Guilford.

- Landerl, K. & Willburger, E. (2009). Der Ein-Minuten-Lese-flüssigkeitstest- ein Verfahren zur Diagnose der Leistung im Wort- und Pseudowortlesen. In W. Lenhard & W. Schneider (Hrsg.), *Diagnostik und Förderung des Leseverständnisses. Tests und Trends. Jahrbuch der pädagogisch- psychologischen Diagnostik* (S. 65–80). Göttingen: Hogrefe.
- Lenhard, W. (2005a). Diagnostische Verfahren zur Schulleistungsfeststellung in der Grundschule. In M. Götz & A. Nießeler (Hrsg.), *Leistung fördern - Förderung leisten* (S. 38–62). Donauwörth: Auer-Verlag.
- Lenhard, W. (2005b). Schwierigkeiten im Schriftspracherwerb. In S. Ellinger & M. Wittrock (Hrsg.), *Sonderpädagogik in der Regelschule: Konzepte-Forschung-Praxis* (S. 257–278). Stuttgart: Kohlhammer
- Lenhard, W. & Artelt, C. (2009). Komponenten des Leseverständnisses. In W. Lenhard & W. Schneider (Hrsg.), *Diagnostik und Förderung des Leseverständnisses. Tests und Trends. Jahrbuch der pädagogisch- psychologischen Diagnostik* (S. 1–17). Göttingen: Hogrefe.
- Lenhard, A. & Lenhard, W. (2009a). *ELFE-Training. Förderung des Leseverständnisses für Schüler der 1. bis 6. Klasse*. Zugriff am 20.05. 2011. Verfügbar unter <http://www.psychometrica.de/elfe-training.html>
- Lenhard, A. & Lenhard, W. (2009b). *Leseverständnis auf Satzniveau*. Zugriff am 16.04.2011. Verfügbar unter http://www.psychometrica.de/elfe1-6_leseverstaendnis2-3.html
- Lenhard, A. & Lenhard, W. (2009c). *Modelle der visuellen Worterkennung*. Zugriff am 14.04.2011. Verfügbar unter http://www.psychometrica.de/elfe1-6_leseverstaendnis2-1.html
- Lenhard, A. & Lenhard, W. (2009d). *Textverständnis*. Zugriff am 16.04.2011. Verfügbar unter http://www.psychometrica.de/elfe1-6_leseverstaendnis2-4.html
- Lenhard, W. & Schneider, W. (2006). *ELFE 1-6: Ein Leseverständnistest für Erst- bis Sechstklässler*. Göttingen: Hogrefe.
- Lenhard W. & Schneider, W. (Hrsg.) (2009). *Diagnostik und Förderung des Leseverständnisses. Tests und Trends. Jahrbuch der pädagogisch-psychologischen Diagnostik*. Göttingen: Hogrefe.

- Mannhaupt, G. (2001). *Lernvoraussetzungen im Schriftspracherwerb. Zur Entwicklung im Vor- und Grundschulalter*. Köln: Kölner Studien.
- Marsh, H. W., Hau, K.-T., Balla, J. R. & Grayson, D. (1998). Is more ever too much? The number of indicators per factor in confirmatory factor analysis. *Multivariate Behavioral Research*, 33 (2), 181–220.
- Marx, H. (1998). *Knuspels Leseaufgaben (KNUSPEL-L). Gruppenlesetest für Kinder Ende des ersten bis vierten Schuljahres*. Göttingen: Hogrefe.
- Marx, H. & Jungmann, T. (2000). Abhängigkeit der Entwicklung des Leseverstehens von Hörverstehen und grundlegenden Lesefertigkeiten im Grundschulalter: Eine Prüfung des Simple View of Reading- Ansatzes. *Zeitschrift für Entwicklungspsychologie und Pädagogische Psychologie*, 32 (2), 81–93.
- McClelland, J. L. & Rumelhart, D. E. (1981). An interactive activation model of context effects in letter perception: Part 1. An account of basic findings. *Psychological Review* 88 (5), 375–407.
- McElvany, N. & Schneider, C. (2009). Förderung von Lesekompetenz. In W. Lenhard & W. Schneider (Hrsg.), *Diagnostik und Förderung des Leseverständnisses. Tests und Trends. Jahrbuch der pädagogisch- psychologischen Diagnostik* (S. 151–183). Göttingen: Hogrefe.
- McKeown, M. G., Beck, I. L., Omanson, R. C. & Perfetti, C. A. (1983). The effects of long-term vocabulary instruction on reading comprehension: A replication. *Journal of Reading Behavior*, 15, 3–18.
- McKoon, G. & Ratcliff, R. (1992). Inference during reading. *Psychological Review*, 99 (3), 440–466.
- Morton, J. (1969). Interaction of information in word recognition. *Psychological Review*, 76 (2), 165–178.
- Morton, J. (1979). Word recognition. In J. Morton & J. C. Marshall (Eds.), *Psycholinguistic series II: Structures and processes* (pp. 107–156). London: Elek Scientific Press.

- National institute of child health & human development (NICHD) (Eds.). (2000). *Reports of the national reading panel: Teaching children to read*. Zugriff am 08.05.2011. Verfügbar unter <http://www.nichd.nih.gov/publications/nrp/findings.cfm>
- Naumann, C. L. (1989). *Gesprochenes Deutsch und Orthographie: Linguistische und didaktische Studien zur Rolle der gesprochenen Sprache in System und Erwerb der Rechtschreibung*. Frankfurt/ Bern/ New York: Peter Lang.
- Oakhill, J. & Garnham, A. (1988). *Becoming a skilled reader*. Oxford: Blackwell.
- OECD (2010). *PISA 2009 results. What students know and can do. Student performance in reading, mathematics and science. Volume I*. Paris: OECD. Zugriff am 16.06.2011. Verfügbar unter http://www.oecd-ilibrary.org/education/pisa-2009-results-what-students-know-and-can-do_9789264091450-en
- Osterhout, L., Holcomb, P. J. & Swinney, D. A. (1994). Brain potentials elicited by garden-path-sentences: Evidence of the application of verb information during parsing. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 20 (4), 786–803.
- Perfetti, C. A. (1985). *Reading ability*. New York: Oxford University Press.
- Perfetti, C. A. (2007). Reading ability: Lexical quality to comprehension. *Scientific Studies of Reading*, 11 (4), 357–383.
- Perfetti, C. A. & Hart, L. (2001). The lexical bases of comprehension skill. In D. S. Gorfien (Ed.), *On the consequences of meaning selection: Perspectives on resolving lexical ambiguity* (pp. 67–86). Washington, DC: American Psychological Association.
- Perfetti, C. A. & Hart, L. (2002). The lexical quality hypothesis. In L. Verhoeven, C. Elbro & P. Reitsma (Eds.), *Precursors of Functional Literacy* (pp. 189–213). Amsterdam/Philadelphia: John Benjamins.
- Perry, C., Ziegler, J.C. & Zorzi, M. (2007). Nested incremental modeling in the development of computational theories: The CDP+ model of reading aloud. *Psychological Review* 114 (2), 273–315.
- Pickering, M. J. & Traxler, M. J. (1998). Plausibility and recovery from garden-paths. An eye tracking study. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 24 (4), 940–961.

- Protopapas, A., Sideridis, G. D., Mouzaki, A. & Simos, P. G. (2007). Development of lexical mediation in the relation between reading comprehension and word reading skills in greek. *Scientific Studies of Reading*, 11 (3), 165–197.
- Rayner, K. (1998). Eye movements in reading and information processing: 20 years of research. *Psychological Bulletin*, 124 (3), 372–422.
- Rayner, K. & Pollatsek, A. (1989). *The psychology of reading*. Englewood Cliffs/N.J: Prentice Hall.
- Reiner, P. (2009). *Strukturgleichungsmodelle in Theorie und Praxis. Zielgruppenanalyse mittels AMOS am Fallbeispiel „Positive Einflussfaktoren auf die Kaufbereitschaft kundenindividueller Produkte“*. Unveröffentlichte Diplomarbeit, Wirtschaftsuniversität, Wien.
- Rost, D. H. & Buch, S. R. (2010). Leseverständnis. In D. H. Rost (Hrsg.), *Handwörterbuch Pädagogische Psychologie* (4., überarbeitete und erweiterte Aufl.) (S. 507–520). Weinheim/Basel: Beltz.
- Rost, D. H. & Hartmann, A. (1992). Lesen, Hören, Verstehen. *Zeitschrift für Psychologie*, 200, 345–361.
- Sassenroth, M. (2003). *Schriftspracherwerb. Entwicklungsverlauf, Diagnostik und Förderung* (5. Aufl.). Bern: Haupt.
- Saxer, U. (1991). Lese(r)forschung-Lese(r)förderung. In A. Fritz (Hrsg.), *Lesen im Medienumfeld* (S. 99–132). Gütersloh: Bertelsmann Stiftung.
- Schabmann, A., Schmidt, B. M., Klicpera, C., Gasteiger-Klicpera, B. & Klingebiel, K. (2009). Does systematic reading instruction impede prediction of reading a shallow orthography? *Psychological Science Quarterly* 51 (2), 315–338.
- Schaffner, E. (2009). Determinanten des Leseverstehens. In W. Lenhard & W. Schneider (Hrsg.), *Diagnostik und Förderung des Leseverständnisses. Tests und Trends. Jahrbuch der pädagogisch-psychologischen Diagnostik* (S. 19–44). Göttingen: Hogrefe.
- Scheerer-Neumann, G. (1997). Lesen und Leseschwierigkeiten. In F.E. Weinert (Hrsg.), *Enzyklopädie der Psychologie, Serie I, Pädagogische Psychologie, Band 3: Psychologie des Unterrichts und der Schule* (S. 279–325). Köln: Kiepenheuer und Witsch.

- Schermelleh-Engel, K., Moosbrugger, H. & Müller, H. (2003). Evaluating the fit of structural equation models: Tests of significance and descriptive goodness-of-fit measures. *Methods of Psychological Research Online*, 8 (2), 23–77. Verfügbar unter <http://www.mpr-online.de>
- Schwantner, U. & Schreiner, C. (Hrsg.). (2010). *PISA 2009. Lesen im elektronischen Zeitalter. Die Ergebnisse im Überblick*. Salzburg: Bundesinstitut für Bildungsforschung, Innovation & Entwicklung des österreichischen Schulwesens.
- Seidenberg, M. S. & McClelland, J. L. (1989). A distributed, developmental model of word recognition and naming. *Psychological Review* 96 (4), 523–568.
- Shankweiler, D. P., Lundquist, E., Dreyer, L. G. & Dickinson, C. C. (1996). Reading and spelling difficulties in high-school students: Causes and consequences. *Reading and writing: An interdisciplinary Journal*, 8, 267–294.
- Shankweiler, D. P., Lundquist, E., Katz, L., Stuebing, K. K., Flechter, J. M., Brady, S. et al. (1999). Comprehension and decoding: Patterns of association in children with reading difficulties. *Scientific Studies of Reading*, 3 (1), 69–94.
- Stahl, S. A. & Fairbanks, M. M. (1986). The effects of vocabulary instruction: A model-based metaanalysis. *Review of Educational Research*, 56 (1), 72–110.
- Stanovich, K. E. (1986). Matthew effect in reading: Some consequences of individual differences in the acquisition of literacy. *Reading Research Quarterly*, 21 (4), 360–407.
- Stanovich, K. E. (1989). Has the learning disabilities field lost its intelligence? *Journal of Learning Disabilities*, 22 (8), 487–492.
- Statistik Austria (2011). *Migration & Integration. Zahlen. Daten. Indikatoren. 2011*. Zugriff am 12.07.2011. Verfügbar unter: http://www.statistik.at/web_de/services/publikationen/2/index.html?id=2&listid=2&detail=621
- Steinbrecher, J. (2007). *Lesesozialisation. Ein Überblick über den Forschungsstand*. Wien: Bundesministerium für Unterricht, Kunst und Kultur.
- Sticht, T. & James, H. J. (1984). Listening and reading. In P. D. Pearson (Ed.), *Handbook of reading research* (pp. 293–317). New York: Longmann.

- Stock, C. (2009). Der Zusammenhang zwischen phonologischer Bewusstheit und Leseleistung bei Grundschulkindern. In W. Lenhard & W. Schneider (Hrsg.), *Diagnostik und Förderung des Leseverständnisses. Tests und Trends. Jahrbuch der pädagogisch-psychologischen Diagnostik* (S. 81–95). Göttingen: Hogrefe.
- Taylor, B. M. & Samuels, S. J. (1983). Children`s use of text structure in the recall of expository material. *American Educational Research Journal*, 20 (4), 517–528.
- Tilstra, J., McMaster, K., van der Broek, P., Kendeou, P. & Rapp, D. (2009). Simple but complex: Components of the simple view of reading across grade levels. *Journal of Research in Reading*, 32 (4), 383–401.
- van Dijk, T. A. & Kintsch, W. (1983). *Strategies of discourse comprehension*. New York: Academic Press.
- Yovanoff, P., Duesbery, L., Alonzo, J., & Tindal, G. (2005). Grade-level invariance of theoretical causal structure predicting reading comprehension with vocabulary and oral reading fluency. *Educational Measurement: Issues and Practice*, 24 (3), 4–12.
- Ziegler, J. C., Perry, C. & Coltheart, M. (2000). The DRC model of visual word recognition and reading aloud: An extension to German. *European Journal of Cognitive Psychology*, 12 (3), 413–430.
- Ziegler, J. C., Perry, C., Jacobs, A. M. & Braun, M. (2001). Identical words are read differently in different languages. *Psychological Science*, 12(5), 379–384.
- Ziegler, J. C., Stone, G. O. & Jacobs, A. M. (1997). What is the pronunciation for -ough and the spelling for /u/? A database for computing feedforward and feedback inconsistency in english. *Behavior Research Methods, Instruments, & Computers*, 29(4), 600–618.
- Zwitserlood, P. (1996). Prozesse und lexikalische Repräsentationen bei der visuellen Worterkennung. In B. Spillner (Hrsg.), *Sprache: Verstehen und Verständlichkeit. Kongressbeiträge der 25. Jahrestagung der Gesellschaft für Angewandte Linguistik, GAL e.V.* (S. 115–118). Frankfurt a. M.: Lang.

11. Abbildungsverzeichnis

Abbildung 1: Schematische Darstellung der Modelle des Leseverständnisses	- 45 -
Abbildung 2: Geschlechtsspezifische Darstellung der Altersverteilung (N=113).....	- 61 -
Abbildung 3: Modell 1 mit standardisierten Lösungen.....	- 66 -
Abbildung 4: Modell 2 mit standardisierten Lösungen.....	- 68 -
Abbildung 5: Modell 3 mit standardisierten Lösungen.....	- 69 -
Abbildung 6: Nested model comparison für Modell 1 und Modell 2 mit standardisierten Lösungen	- 71 -

12. Tabellenverzeichnis

Tabelle 1: Deskriptive Beschreibung der Testergebnisse für die jeweilige Variable der Stichprobe (N = 113).....	- 64 -
Tabelle 2: Korrelationen im Modellvergleich	- 64 -
Tabelle 3: Vergleich der drei Strukturgleichungsmodelle.....	- 70 -
Tabelle 4: Koeffizienten, Critical Ratio (C.R.) und Signifikanz (p) der direkten Pfade im präferierten Modell 1	- 72 -

13. Anhang

Anhang A: Ansuchen zur Studienbewilligung an den Landesschulrat

28.06.2010

Ao. Univ. Prof. Dr. Mag. Alfred Schabmann
Universität Wien
Institut für Wirtschaftspsychologie, Bildungspsychologie und Evaluation
Universitätsstraße 7
A-1010 Wien
Tel.: (+1) 4277 478 92
e-mail: alfred.schabmann@univie.ac.at

An den
Landesschulrat für Oberösterreich

c/o Fritz Enzenhofer
Sonnensteinstraße 20
A-4040 Linz
0732/7071-0

Betrifft: Ansuchen um Genehmigung einer wissenschaftlichen Studie der Universität Wien zum Thema „*Die Mediatorrolle des Wortschatzes in Bezug auf das Leseverständnis bei Grundschulkindern*“

Sehr geehrter Herr Enzenhofer,

An der Fakultät für Psychologie der Universität Wien wird zurzeit ein Forschungsprojekt zum Thema „*Die Mediatorrolle des Wortschatzes in Bezug auf das Leseverständnis bei Grundschulkindern*“ durchgeführt. Dabei geht es im Wesentlichen darum zu erfahren, wie sich der Wortschatz eines Kindes auf das Leseverständnis auswirkt.

Für diese Untersuchung planen wir etwa jeweils 100–120 Kinder der zweiten, dritten und vierten Schulstufe auf ihr Leseverständnis zu testen. Konkret werden den Kindern mehrere verschiedene Aufgaben zum Lesen vorgegeben. Zum einen werden die Dekodierfähigkeiten Leseflüssigkeit und Lesegenauigkeit getestet. Ebenso werden das Leseverständnis, das mündliche sprachliche Verstehen sowie der Wortschatz der Kinder getestet. Die Erhebung ist für September/Oktober 2010 geplant und wird pro Kind etwa eine halbe Stunde benötigen.

Für die Untersuchung haben wir vier Schulen kontaktiert und auch bereits die Zusage der DirektorInnen erhalten. Es handelt sich dabei um die VS Schörfling (Gmundnerstraße 21, 4861 Schörfling), die VS Mondsee (Schulweg 4, 5310 Mondsee), VS St. Lorenz – Tiefgraben (Thalgaustraße 4, 5310 Mondsee) VS Seewalchen (Schulweg, 4863 Seewalchen). Es ist selbstverständlich, dass die bei der Erhebung anfallenden Daten vertraulich behandelt werden und auch das Einverständnis der Eltern eingeholt wird. Wir bitten um Genehmigung der Untersuchung. Wir bitten Sie um Rückantwort an folgende Adresse:

Alfred Schabmann
(Projektleitung),
Stefanie Dorn, Yvonne
Huemer, Maria Klausecker

Mit freundlichen Grüßen,

Anhang B: Bewilligung der Studie durch den Landesschulrat

LANDESSCHULRAT FÜR OBERÖSTERREICH



A-4040 LINZ, SONNENSTEINSTRASSE 20

Herrn Univ.-Prof. Mag. Dr. Alfred Schabmann
Universität Wien
Institut für Wirtschaftspsychologie
Universitätsstraße 7
1010 Wien

Bearbeiterin:
U. Wagner
Tel: 0732/7071-2321
Fax: 0732/7071-2330
E-mail:lsr@lsr-ooe.gv.at

Ihr Zeichen	Vom	Unser Zeichen	Vom
---	28. 6. 2010	B5 – 14/42 – 2010	2. 7. 2010

**Untersuchung im Rahmen der wissenschaftlichen Studie
"Die Mediatorrolle des Wortschatzes in Bezug auf das
Leseverständnis bei Grundschulkindern"**

Sehr geehrte....

Ihr Ansuchen an Herrn Präsident Enzenhofer um Genehmigung der gegenständlichen Untersuchung an oö. Volksschulen wurde zur Bearbeitung an die zuständige Abteilung Schulpsychologie-Bildungsberatung weitergeleitet.

Nach Prüfung der Unterlagen genehmigt der Landesschulrat für OÖ Ihre Untersuchung unter den üblichen Bedingungen:

Freiwilligkeit der Teilnahme
Information der Erziehungsberechtigten
Einhaltung der Datenschutzbestimmungen
Übermittlung des Endberichtes an den Landesschulrat für OÖ.

Bei Ihrer Kontaktaufnahme mit den Schulen verweisen Sie bitte auf diese Genehmigung.

Mit freundlichen Grüßen

Für den Amtsführenden Präsidenten
Dr. Lang eh.

Anhang C: Elternbrief

Liebe Eltern!

Im Rahmen eines Forschungsprojektes an der Fakultät für Psychologie der Universität Wien wird eine Studie zum Leseverständnis bei Grundschulkindern durchgeführt.

Es werden ein Lesetest (lautes „Wortlesen“), ein Wortschatztest, ein Leseverständnistest und ein Test zur Erfassung des Hörverstehens vorgegeben. Die Testungen werden möglichst kurz gehalten und im Zuge einer Gruppentestung der gesamten Klasse, sowie in einer Einzeltestung durchgeführt.

Wir versichern, dass die Auswertung der Ergebnisse vertraulich erfolgt. Speziell die Lehrer/innen ihrer Kinder werden nicht über die Leistungen informiert. Wir sind zur Verschwiegenheit und anonymen Verarbeitung der Daten verpflichtet.

Herzlichen Dank, dass Sie eine wissenschaftliche Studie zu diesem Thema möglich machen!

Bei Rückfragen stehe ich ihnen gerne zur Verfügung:

Maria Klausecker

Gmundnerstraße 27

4861 Schörfling

Tel.Nr.: 0680 1185533

m.klausecker@gmail.com

Mit freundlichen Grüßen,

Prof. Dr. Mag. Alfred Schabmann (Projektleitung)

Stefanie Dorn (Diplomandin)

Yvonne Huemer (Diplomandin)

Maria Klausecker (Diplomandin)

Name: _____

Schule/Klasse: _____

- ☐ Ja, ich bin damit einverstanden, dass mein Kind an der Studie teilnimmt.
- ☐ Nein, ich bin nicht einverstanden, dass mein Kind an der Studie teilnimmt.

Unterschrift: _____

Anhang D: Testvorgabe

Deckblatt (demographische Daten)

Liebe Kinder!

Wir werden heute gemeinsam einige Übungen zum Thema lesen machen. Bei den Übungen sollt ihr euch anstrengen und sie alleine machen.

Wenn ihr eine Übung einmal nicht schafft, macht das gar nichts. Diese Aufgaben sind dann schon für ältere Kinder bestimmt.



Zu Beginn sollt ihr bitte schon einmal die erste Seite ausfüllen. Diese Dinge haben noch nichts mit den Übungen zu tun.

Vorname: _____

Nachname: _____

Klasse: _____

Schule: _____

Mädchen: ☐

Bub: ☐

Geburtsdatum (Tag, Monat, Jahr): _____

Alter: ____ Jahre

Wohnort: _____

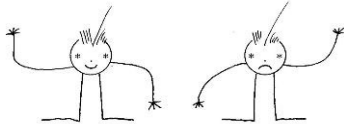
Muttersprache: deutsch ☐ andere ☐

Wochentag: _____

Knuspel-L – Untertest Hörverstehen (Marx, 1998)

Knuspels Leseaufgaben

Form A



A

☞

B

☞

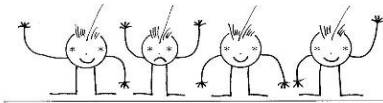
Junge

Mädchen

Bitte warten!

© by Hegner-Verlag GmbH & Co. KG, Göttingen - Nachdruck und jegliche Art der Vervielfältigung verboten - Best.-Nr. 974/07/95

2



1|

☞

2|

☞

glatt

lockig

3|

☞

keine

1.

2.

3.

4.

5.

4|

☞



5|

☞

6

7

8

9

10

11

12

13

14

6|

☞

knuspern



7|

☞

Bitte weitermachen auf Seite 3 oben -->



8|

☞

9|

☞

10|

☞



11|

☞

Montag Dienstag Mittwoch Donnerstag Freitag Samstag Sonntag

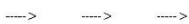
12|

☞



13|

☞



14|

☞

auf dem Bauch

auf den Beinen

Bitte warten!

3

ELFE 1-6 – Untertest Textverständnis (Lenhard & Schneider, 2006)

Du siehst hier kleine Geschichten mit einer Frage.
Bitte streiche die richtige Antwort an!

Beispiele:

Heute scheint den ganzen Tag die Sonne.

Welcher Satz stimmt?

- ☐ Heute ist schönes Wetter. ☐ Morgen wird es regnen.
☐ Gestern war schönes Wetter. ☐ Heute regnet es.

Ein Pferd, das ist ein großes Tier. Es hat auch Beine und zwar vier.

Ein Pferd...

- ☐ ist ganz klein. ☐ hat vier Beine.
☐ hat braune Haare. ☐ frisst Gras.



Stopp! Noch nicht umblättern!

11

Tim freut sich, wenn die Sonne scheint. Dann kann er mit seinen Freunden Fußball spielen.

Tim....

- ☐ isst gerne Obst. ☐ macht seine Hausaufgaben.
☐ ärgert seine Schwester. ☐ spielt gerne Fußball.

1

Felix spielt mit seinem schönen neuen Ball. Felix sagt, dass Jan und Eva nicht mitspielen dürfen. Deshalb sind sie böse auf ihn.

Felix spielt...

- ☐ nicht mit Jan und Eva. ☐ mit dem Ball von Jan und Eva.
☐ mit Jan und Eva. ☐ gern mit Jan und Eva.

2

Jan und Eva sind böse, weil...

- ☐ er einen neuen Ball hat. ☐ er sie nicht mitspielen lässt.
☐ sie nicht gern spielen. ☐ er nicht mit dem Ball spielt.

3

Evi und ihr großer Bruder Stefan wollen fernsehen. Sie können sich aber nicht auf ein Programm einigen und fangen an zu streiten. Stefan nimmt Evi die Fernbedienung weg und schaltet auf seine Lieblingssendung. Evi sagt: „Das ist gemein! Immer machst du was du willst, nur weil du der Stärkere bist!“

Welcher Satz ist richtig?

- ☐ Evi ist stärker als ihr großer Bruder. ☐ Evi möchte etwas anderes anschauen als Stefan.
☐ Stefan und Evi möchten dieselbe Sendung anschauen. ☐ Evi und Stefan streiten sich nie.

4

Paula ist mit ihren Eltern in den Ferien ans Meer gefahren. Am Strand spielt sie im Sand und sammelt schöne, farbige Muscheln. Die findet sie so schön.

Paula...

- ☐ ist mit ihren Eltern in die Berge gefahren. ☐ schwimmt gerne im Meer.
☐ hat Angst vor Krebsen. ☐ mag farbige Muscheln.

5

12

Wale legen keine Eier sondern bringen ihre Jungen lebend zur Welt. Die Jungen trinken bei ihrer Mutter Milch. Deshalb sind Wale keine Fische, sondern Säugetiere.

Wale....

- ☐ sind Fische. ☐ sind Säugetiere.
☐ legen Eier. ☐ fressen am liebsten Fische.

6

Wale leben im Meer. Zum Atmen schwimmen sie zur Wasseroberfläche. Dort holen sie tief Luft und können dann lange unter Wasser tauchen.

Wale....

- ☐ können unter Wasser atmen. ☐ leben in Seen.
☐ müssen zum Luftholen zur Wasseroberfläche kommen. ☐ können nur kurz tauchen.

7

Lars muss für seine Mutter einkaufen. Im Laden kann er aber das Geld nicht finden. Hat er es verloren? Mutter wird böse sein. Er erzählt der Mutter aus Angst eine erfundene Geschichte über einen Dieb. Die Mutter schüttelt den Kopf und sagt: „Du hast das Geld hier vergessen!“ Da wird Lars rot und schämt sich sehr.

Lars erfindet eine Ausrede, weil....

- ☐ er glaubt, das Geld verloren zu haben. ☐ er nicht einkaufen will.
☐ ein Dieb ihm das Geld gestohlen hat. ☐ die Mutter den Kopf schüttelt.

8

Mutter weiß, dass....

- ☐ Lars die Wahrheit sagt. ☐ Lars das Geld verloren hat.
☐ ein Dieb ihm das Geld gestohlen hat. ☐ Lars geschwindelt hat.

9

Lars sagt zu seiner Mutter, dass....

- ☐ er das Geld verloren hat. ☐ ein Dieb ihm das Geld gestohlen hat.
☐ er das Geld im Laden gefunden hat. ☐ er Angst hat.

10

Die Kinder spielen verstecken. Fast jeder hat ein gutes Versteck. Alex ist leicht zu finden.

Alex....

- ☐ findet die anderen Kinder leicht. ☐ hat ein gutes Versteck.
☐ ist der Fänger. ☐ hat ein schlechtes Versteck.

11

13

Anna und Martin dürfen heute nicht draußen spielen. Es ist zu kalt. Sie helfen der Mutter beim Kuchen backen. „Das Backen macht zwar keinen Spaß“, sagt Martin zu Anna, „aber der Kuchen schmeckt gut.“

Die Kinder dürfen heute....

- ☐ der Mutter nicht in der Küche helfen. ☐ draußen in der Kälte spielen.
☐ nur im Haus spielen. ☐ keinen Kuchen essen.

12

Martin ist gerne Kuchen, aber....

- ☐ das Backen macht ihm keinen Spaß. ☐ die Mutter sagt, dass er keinen Kuchen essen darf.
☐ er hilft gern beim Backen. ☐ Anna sagt, dass Backen keinen Spaß macht.

13

Was wird in dieser Geschichte erzählt?

- ☐ Martin will Kuchen backen, weil es draußen zu kalt ist. ☐ Die Kinder backen Kuchen, anstatt draußen in der Kälte zu spielen.
☐ Die Kinder spielen zuerst draußen, dann helfen sie beim Kuchen backen. ☐ Martin sagt zu Anna, dass er lieber bäckt als spielt.

14

Nicki ist der einzige Hase mit kurzen Ohren. Alle anderen Hasen lachen ihn deshalb aus. Aber Nicki lacht auch, denn er weiß, dass Jäger lange Ohren besser sehen können als kurze Ohren.

Nicki....

- ☐ hat lange Ohren. ☐ ist ein Jäger.
☐ hat keine langen Ohren. ☐ ist kein Hase.

15

Nicki lacht, weil er weiß, dass Jäger....

- ☐ kurze Ohren besser sehen können. ☐ kurze und lange Ohren gleich gut sehen können.
☐ Hasen nicht sehen können. ☐ die Ohren der anderen Hasen besser sehen können.

16

14

Tina muss heute als Hausaufgabe eine Geschichte lesen. Sie hat keine Lust dazu. Endlich fängt sie an. Es ist eine spannende Geschichte. Tina staunt: Hausaufgaben können auch Spaß machen.

Tina...

- | | |
|---|--|
| ☐ liest gern, aber die Geschichte hat ihr nicht gefallen. | ☐ hatte zuerst keine Lust zu lesen, aber die Geschichte hat ihr dann gefallen. |
| ☐ hat die Geschichte gelesen, weil sie keine Lust hatte. | ☐ hat vergessen, ihre Hausaufgaben zu machen. |

17

Tina hat...

- | | |
|--|---|
| ☐ eine langweilige Geschichte gelesen. | ☐ eine spannende Geschichte geschrieen. |
| ☐ ihre Hausaufgaben nicht gemacht. | ☐ etwas Spannendes gelesen. |

18

Vor vielen tausend Jahren lebten in Europa große behaarte Elefanten, die Mammuts. Gegen Ende der Eiszeit starben diese Tiere jedoch aus. Man weiß heute sehr genau wie sie aussahen, weil man einige Mammuts im Dauerfrostboden Sibiriens gefunden hat. Dort waren sie wie in einer Gefriertruhe eingefroren.

Was steht im Text?

- | | |
|--|---|
| ☐ Mammuts wurden von den Steinzeitmenschen gejagt. | ☐ Mammuts hatten keine Haare. |
| ☐ Einige Mammuts sind seit der Eiszeit im Boden Sibiriens eingefroren. | ☐ Mammuts hatten eine dicke Speckschicht. |

19

Lena ist die beste Freundin von Steffi. Sie wollen heute nach der Schule zusammen spielen. Steffi hat Lena versprochen zu kommen.

Wer kommt zu wem?

- | | |
|---|---|
| ☐ Lena kommt zum Spielen zu Steffi. | ☐ Die beiden Mädchen treffen sich auf dem Spielplatz. |
| ☐ Steffi kommt zum Spielen zu Lena. | ☐ Jeder bleibt heute daheim. |

20



Stopp! Hier ist der Test zu Ende!

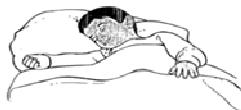
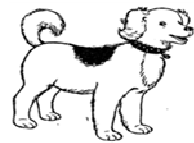
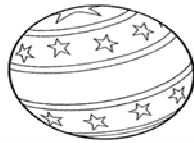
adaptierte Version des PPVT-III (Dunn & Dunn, 1997)

Set 5, 7, 8, 10, 12

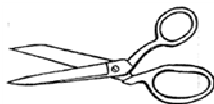
PPVT Übungsaufgaben



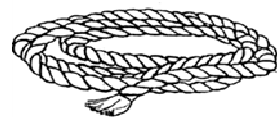
A



B



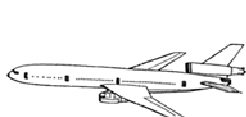
C

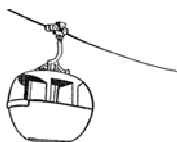


D

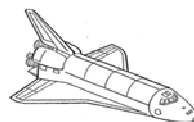


13. Anhang







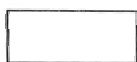




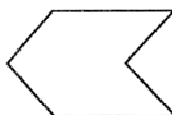


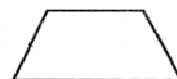










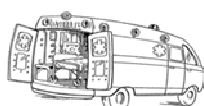










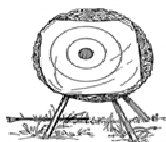


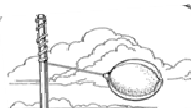


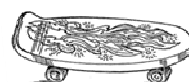




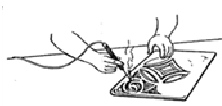


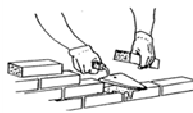


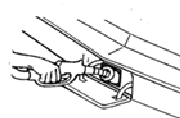




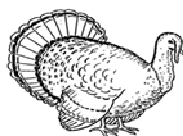






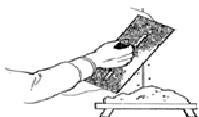














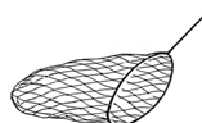




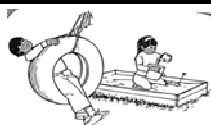






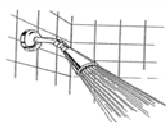


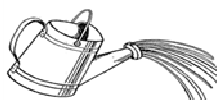








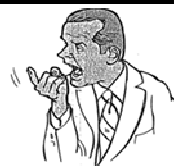
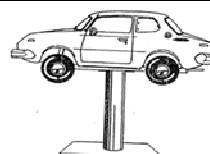
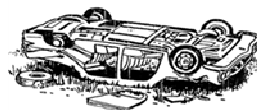
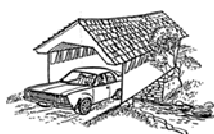
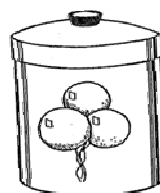
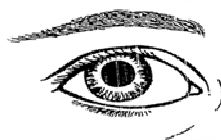


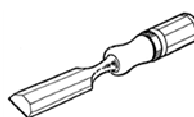
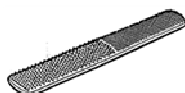
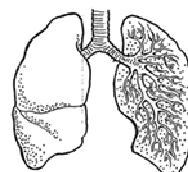
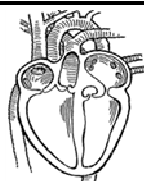
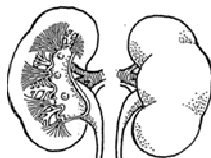




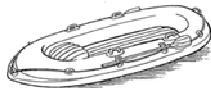
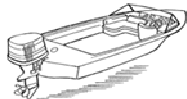
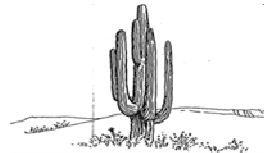
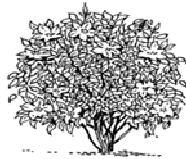
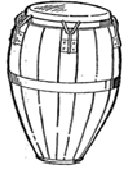
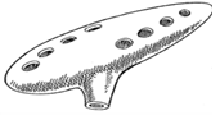


13. Anhang

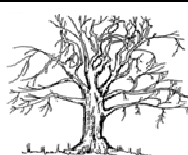
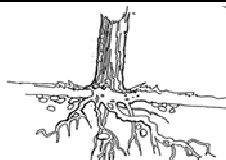
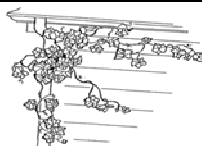
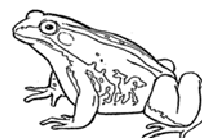
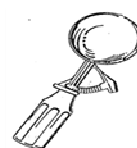
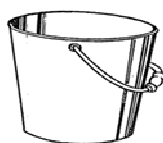
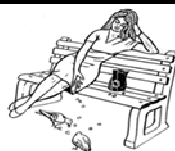
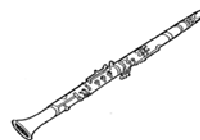
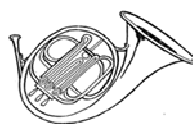
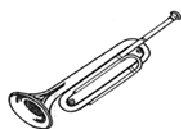




13. Anhang



13. Anhang



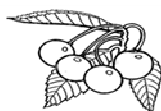
13. Anhang







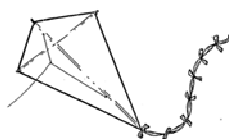


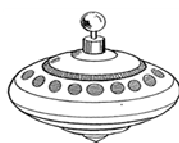






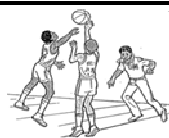








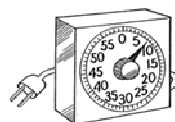


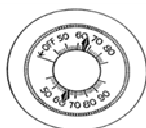


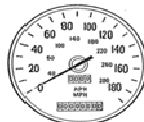


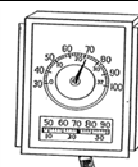












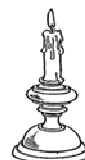
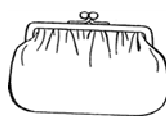
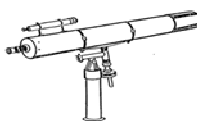
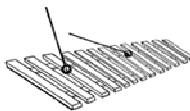




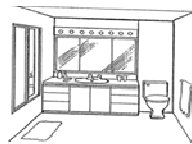
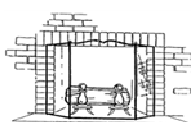
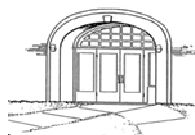
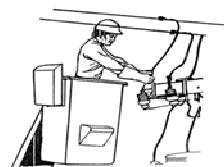
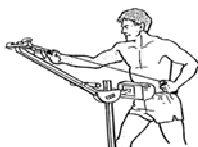
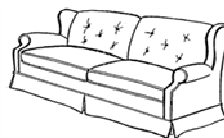
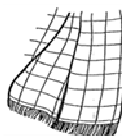
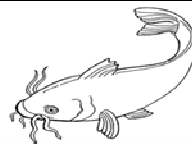
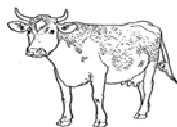


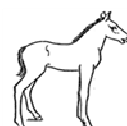
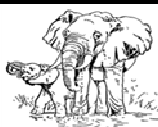
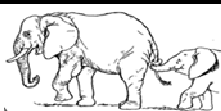
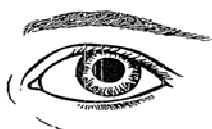
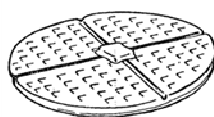
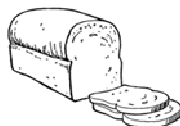
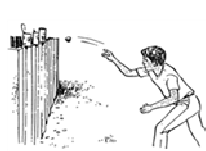


13. Anhang



13. Anhang





Wortlesetest (Schabmann, Schmidt, Klicpera, Gasteiger-Klicpera & Klingebiel, 2009)

Testvorgabe erfolgte über einzelne Wörter in einer Power-Point-Präsentation

NAME:.....

Form A

KLASSE:.....

SCHULE:.....

Haus	Tee	Wind	
Tag	Dach	Hund	
Post	Geld	Haut	Sek:.....
Bad	Mai	Topf	Fehler:.....
Brot	Tisch	Ziel	SK:.....

Seil	Gras	Dieb	
Zoo	Kuh	Spatz	
Knie	Ton	Gas	Sek:.....
Bär	Kreis	Nacht	Fehler:.....
Saal	Schmutz	Land	SK:.....

Dezember	Erziehung	Minute	
Ausbildung	Februar	Fernseher	
Theater	Geheimnis	Oktober	Sek:.....
Geburtstag	Polizei	Fahrerin	Fehler:.....
Ferien	Beamtin	Kartoffel	SK:.....

Rakete	Widerspruch	Alkohol	
Abteilung	Lineal	Heiterkeit	
Konferenz	Darstellung	Toilette	Sek:.....
Omnibus	Silvester	Einschreiben	Fehler:.....
Direktor	Zustimmung	Kamerad	SK:.....

Moch	Faka	Splant	
Kaus	Fans	Breigt	
Kerd	Meil	Pischt	Sek:.....
Mart	Zang	Frilp	Fehler:.....
Geus	Plit	Spein	SK:.....
Frunkelte	Lintabel	Agelos	
Kolehner	Verlanen	Feltliche	
Spieroge	Piralet	Mandriche	Sek:.....
Parender	Verbalut	Schelperta	Fehler:.....
Vertucke	Getusin	Sprutzende	SK:.....

Anhang E: Ergebnisse der Strukturgleichungsmodelle (Amos-Output)*Modell 1***Notes for Model (Default model)**

Computation of degrees of freedom (Default model)

Number of distinct sample moments: 15

Number of distinct parameters to be estimated: 11

Degrees of freedom (15 - 11): 4

Result (Default model)

Minimum was achieved

Chi-square = 2,827

Degrees of freedom = 4

Probability level = ,587

Model Fit Summary

CMIN

Model	NPAR	CMIN	DF	P	CMIN/DF
Default model	11	2,827	4	,587	,707
Saturated model	15	,000	0		
Independence model	5	103,909	10	,000	10,391

Model	RMR	GFI	AGFI	PGFI
Default model	2,016	,990	,963	,264
Saturated model	,000	1,000		
Independence model	17,600	,735	,603	,490

Baseline Comparisons

Model	NFI Delta1	RFI rho1	IFI Delta2	TLI rho2	CFI
Default model	,973	,932	1,012	1,031	1,000
Saturated model	1,000		1,000		1,000

Model	NFI Delta1	RFI rho1	IFI Delta2	TLI rho2	CFI
Independence model	,000	,000	,000	,000	,000

RMSEA

Model	RMSEA	LO 90	HI 90	PCLOSE
Default model	,000	,000	,122	,705
Independence model	,290	,241	,341	,000

Maximum Likelihood Estimates**Regression Weights: (Group number 1 - Default model)**

		Estimate	S.E.	C.R.	P	Label
LT_Summe_Fehler <---	Worterkennen	1,000				
LT_Summe_Zeit <---	Worterkennen	12,726	4,071	3,126	,002	par_1
ELFE_gesamt <---	Worterkennen	-1,329	,414	-3,212	,001	par_2
ELFE_gesamt <---	PPVT_gesamt	,168	,089	1,891	,059	par_3
ELFE_gesamt <---	Knuspiel_gesamt	,486	,116	4,189	***	par_4

Standardized Regression Weights: (Group number 1 - Default model)

		Estimate
LT_Summe_Fehler <---	Worterkennen	,400
LT_Summe_Zeit <---	Worterkennen	,753
ELFE_gesamt <---	Worterkennen	-,657
ELFE_gesamt <---	PPVT_gesamt	,150
ELFE_gesamt <---	Knuspiel_gesamt	,332

Matrices (Group number 1 - Default model)**Residual Covariances (Group number 1 - Default model)**

	Knuspiel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler
Knuspiel_gesamt	,000				
PPVT_gesamt	,000	,000			
ELFE_gesamt	,240	-,363	,111		
LT_Summe_Zeit	-1,837	7,395	,348	,000	
LT_Summe_Fehler	-1,374	-,411	-,736	,000	,000

Standardized Residual Covariances (Group number 1 - Default model)

	Knuspel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler
Knuspel_gesamt	,000				
PPVT_gesamt	,000	,000			
ELFE_gesamt	,263	-,315	,063		
LT_Summe_Zeit	-,260	,802	,030	,000	
LT_Summe_Fehler	-1,315	-,301	-,467	,000	,000

Modell 2**Notes for Model (Default model)**

Computation of degrees of freedom (Default model)

Number of distinct sample moments: 15

Number of distinct parameters to be estimated: 12

Degrees of freedom (15 - 12): 3

Result (Default model)

Minimum was achieved

Chi-square = 2,404

Degrees of freedom = 3

Probability level = ,493

Model	NPAR	CMIN	DF	P	CMIN/DF
Default model	12	2,404	3	,493	,801
Saturated model	15	,000	0		
Independence model	5	103,909	10	,000	10,391

Model	RMR	GFI	AGFI	PGFI
Default model	1,257	,992	,958	,198
Saturated model	,000	1,000		
Independence model	17,600	,735	,603	,490

Model	NFI Delta1	RFI rho1	IFI Delta2	TLI rho2	CFI
Default model	,977	,923	1,006	1,021	1,000
Saturated model	1,000		1,000		1,000
Independence model	,000	,000	,000	,000	,000

Model	RMSEA	LO 90	HI 90	PCLOSE
Default model	,000	,000	,147	,605
Independence model	,290	,241	,341	,000

	Estimate	S.E.	C.R.	P	Label
--	----------	------	------	---	-------

			Estimate	S.E.	C.R.	P	Label
Worterkennen	<---	PPVT_gesamt	,043	,067	,643	,520	par_6
LT_Summe_Fehler	<---	Worterkennen	1,000				
LT_Summe_Zeit	<---	Worterkennen	12,959	4,191	3,092	,002	par_1
ELFE_gesamt	<---	Worterkennen	-1,335	,418	-3,198	,001	par_2
ELFE_gesamt	<---	PPVT_gesamt	,191	,097	1,970	,049	par_3
ELFE_gesamt	<---	Knuspel_gesamt	,486	,116	4,195	***	par_4
			Estimate				
Worterkennen	<---	PPVT_gesamt	,079				
LT_Summe_Fehler	<---	Worterkennen	,396				
LT_Summe_Zeit	<---	Worterkennen	,760				
ELFE_gesamt	<---	Worterkennen	-,660				
ELFE_gesamt	<---	PPVT_gesamt	,172				
ELFE_gesamt	<---	Knuspel_gesamt	,336				
	Knuspel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler		
Knuspel_gesamt	,000						
PPVT_gesamt	,000	,000					
ELFE_gesamt	,361	,000	,351				
LT_Summe_Zeit	-3,854	1,526	-1,501	,000			
LT_Summe_Fehler	-1,529	-,864	-,955	,000	,000		
	Knuspel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler		
Knuspel_gesamt	,000						
PPVT_gesamt	,000	,000					
ELFE_gesamt	,401	,000	,205				
LT_Summe_Zeit	-,546	,165	-,132	,000			
LT_Summe_Fehler	-1,464	-,633	-,613	,000	,000		

Modell 3**Notes for Model (Default model)**

Computation of degrees of freedom (Default model)

Number of distinct sample moments: 15

Number of distinct parameters to be estimated: 12

Degrees of freedom (15 - 12): 3

Result (Default model)
 Minimum was achieved
 Chi-square = 1,813
 Degrees of freedom = 3
 Probability level = ,612

Model	NPAR	CMIN	DF	P	CMIN/DF
Default model	12	1,813	3	,612	,604
Saturated model	15	,000	0		
Independence model	5	103,909	10	,000	10,391
Model	RMR	GFI	AGFI	PGFI	
Default model	,758	,994	,968	,199	
Saturated model	,000	1,000			
Independence model	17,600	,735	,603	,490	
Model	NFI	RFI	IFI	TLI	CFI
	Delta1	rho1	Delta2	rho2	
Default model	,983	,942	1,012	1,042	1,000
Saturated model	1,000		1,000		1,000
Independence model	,000	,000	,000	,000	,000
Model	RMSEA	LO 90	HI 90	PCLOSE	
Default model	,000	,000	,131	,709	
Independence model	,290	,241	,341	,000	

		Estimate	S.E.	C.R.	P	Label
PPVT_gesamt	<--- Worterkennen	,199	,204	,973	,330	par_5
LT_Summe_Fehler	<--- Worterkennen	1,000				
LT_Summe_Zeit	<--- Worterkennen	12,831	4,103	3,128	,002	par_1
ELFE_gesamt	<--- Worterkennen	-1,339	,420	-3,191	,001	par_2
ELFE_gesamt	<--- PPVT_gesamt	,208	,102	2,049	,040	par_3
ELFE_gesamt	<--- Knuspel_gesamt	,462	,120	3,854	***	par_4

		Estimate
PPVT_gesamt	<--- Worterkennen	,109
LT_Summe_Fehler	<--- Worterkennen	,398
LT_Summe_Zeit	<--- Worterkennen	,756
ELFE_gesamt	<--- Worterkennen	-,663
ELFE_gesamt	<--- PPVT_gesamt	,188
ELFE_gesamt	<--- Knuspel_gesamt	,318

	Knuspel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler
Knuspel_gesamt	,000				
PPVT_gesamt	-,035	-,042			
ELFE_gesamt	,231	,116	,271		
LT_Summe_Zeit	-1,837	-,678	-,952	,000	

	Knuspel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler
LT_Summe_Fehler	-1,374	-1,040	-,872	,000	,000
	Knuspel_gesamt	PPVT_gesamt	ELFE_gesamt	LT_Summe_Zeit	LT_Summe_Fehler
Knuspel_gesamt	,000				
PPVT_gesamt	-,043	-,030			
ELFE_gesamt	,255	,102	,157		
LT_Summe_Zeit	-,260	-,073	-,084	,000	
LT_Summe_Fehler	-1,315	-,760	-,557	,000	,000

Modellvergleich

Nested model comparison für Modell 1 und Modell 2

Nested Model Comparisons

Assuming model mit reg to be correct:

Model	DF	CMIN	P	NFI Delta-1	IFI Delta-2	RFI rho-1	TLI rho2
ohne reg	1	,423	,515	,004	,004	-,009	-,010

Nested model comparison für Modell 1 und Modell 3

Nested Model Comparisons

Assuming model mit reg to be correct:

Model	DF	CMIN	P	NFI Delta-1	IFI Delta-2	RFI rho-1	TLI rho2
ohne reg	1	1,014	,314	,010	,010	,010	,011

Anhang F: Curriculum vitae

Persönliche Daten

Name: Stefanie Dorn

Geburtsdatum: 23.04.1985

Geburtsort: Würzburg, Deutschland

Staatsbürgerschaft: deutsch

Anschrift: Hernalser Gürtel 4/15
1080 Wien

E-Mail: stefanie.dorn@hotmail.com

Bildung

Seit 10/2005 Diplomstudium Psychologie an der Universität Wien
mit Schwerpunktausbildung Klinische Psychologie

03/2009 – 01/2010 universitäres Ausbildungsprogramm zum student mentor

07/2009 – 09/2009 Fachpraktikum an der Medizinischen Universität Wien,
Besonderer Einrichtung für Medizinische Aus- und
Weiterbildung, Abteilung Personalentwicklung

Aufgaben: Mitarbeit bei Entwicklung, Planung,
Organisation von Personalentwicklungsmaßnahmen und
Konzeptentwicklung

09/1995 – 06/2004 Städtisches Mozart-Schönborn Gymnasium Würzburg
(G9), Abschluss: Abitur (Notendurchschnitt 1,9)

09/1991 – 07/1995 Grundschule Bechtolsheimer Hof, Würzburg

Berufliche Tätigkeit

2008 – 2011	geringfügige Beschäftigungen u.a. Mitarbeit bei Bettschart & Kofler, Medien- und Kommunikationsberatung GmbH (z. B. Frauengesundheitstage, Tag der seelischen Gesundheit im Wiener Rathaus)
03/2005 – 09/2005	Aushilfstätigkeit als Pflegehilfskraft beim Arbeiter-Samariter-Bund Würzburg e.V.
08/2004 – 02/2005	Pflegehilfskraft in der Abteilung Individuellen Schwerstbehindertenpflege des Arbeiter-Samariter-Bundes Würzburg e.V. Aufgaben: Persönliche Assistenz (24 Std. Pflege)

Sprachkenntnisse

Deutsch:	Muttersprache
Englisch:	fließend in Wort und Schrift
Französisch:	gesprochen fließend; gute Kenntnisse schriftlich