

Orthographieerwerb im Kontext Deutsch als Fremdsprache

Eine korpusbasierte Untersuchung zur
Rechtschreibkompetenz chinesischer
Deutschlernender

Liu Ming

Thema Sprache –
Wissenschaft für den Unterricht



| wbv

Orthographieerwerb im Kontext Deutsch als Fremdsprache

Eine korpusbasierte Untersuchung zur
Rechtschreibkompetenz chinesischer Deutschlernender

Liu Ming

Reihe „Thema Sprache – Wissenschaft für den Unterricht“

Ziel dieser Reihe ist die Auslotung des wissenschaftlichen Potenzials, das eine Beschäftigung mit Sprache in Bezug auf schulische Kontexte hat. Dabei wird zum einen gefragt, wie Anwendungsfelder und Erkenntnisse der wissenschaftlichen Disziplinen Sprachwissenschaft, Sprachlehrforschung und den Sprachdidaktiken gewinnbringend für schulische Ziele in die Unterrichtspraxis übertragen werden können, zum anderen, welche Impulse aus dem Unterricht für die wissenschaftlichen Disziplinen ausgehen könnten.

Die Reihe unterliegt dem doppelt blinden Peer-Reviewverfahren.

Die Reihe wird herausgegeben von **Anja Binanzer** (Technische Universität Dresden), **Miriam Langlotz** (Universität Kassel), **Verena Wecker** (Westfälische Wilhelms-Universität Münster), **Maurice Fürstenberg** (Ludwig-Maximilians-Universität München) und **Hans-Georg Müller** (Universität Potsdam).

Wissenschaftlicher Beirat:

Ursula Bredel (Hildesheim), Doreen Bryant (Tübingen), Nicole Marx (Köln), Anja Müller (Mainz), Iris Rautenberg (Ludwigsburg), Claudia Riemer (Bielefeld), Michael Rödel (München), Björn Rothstein (Bochum), Rosemarie Tracy (Mannheim), Constanze Weth (Luxembourg)

Liu Ming

Orthographieerwerb im Kontext Deutsch als Fremdsprache

**Eine korpusbasierte Untersuchung zur
Rechtschreibkompetenz chinesischer
Deutschlernender**



| **wbv**

Originaltitel der Dissertation: „Orthographiekompetenz chinesischer Deutschlernender – eine korpusbasierte Untersuchung“

Ein Schneider-Titel bei
wbv Publikation
2026 wbv Publikation
ein Geschäftsbereich der
wbv Media GmbH & Co. KG, Bielefeld

Gesamtherstellung:
wbv Media GmbH & Co. KG
Auf dem Esch 4, 33619 Bielefeld,
service@wbv.de
wbv.de

Umschlaggestaltung:
Christiane Zay, Passau

ISBN Print: 978-3-7639-8034-5
ISBN E-Book: 978-3-7639-8035-2
DOI: 10.3278/9783763980352

Printed in Germany

Diese Publikation ist frei verfügbar zum Download unter
www.wbv.de/ISBN/9783763980352

Diese Publikation ist unter folgender Creative-Commons-
Lizenz veröffentlicht:
creativecommons.org/licenses/by-sa/4.0/deed.de



Für alle in diesem Werk verwendeten Warennamen sowie Firmen- und Markenbezeichnungen können Schutzrechte bestehen, auch wenn diese nicht als solche gekennzeichnet sind. Deren Verwendung in diesem Werk berechtigt nicht zu der Annahme, dass diese frei verfügbar seien.

Der Verlag behält sich das Text- und Data-Mining nach § 44b UrhG vor, was hiermit Dritten ohne Zustimmung des Verlages untersagt ist.

Bibliografische Information der Deutschen Nationalbibliothek

Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliografie; detaillierte bibliografische Daten sind im Internet über <http://dnb.d-nb.de> abrufbar.

Die freie Verfügbarkeit der E-Book-Ausgabe dieser Publikation wurde ermöglicht durch ein Netzwerk wissenschaftlicher Bibliotheken und Institutionen zur Förderung von Open Access in den Sozial- und Geisteswissenschaften im Rahmen der *wbv OpenLibrary 2025*.

Die Publikation beachtet unsere Qualitätsstandards für Open-Access-Publikationen, die an folgender Stelle nachzulesen sind:

https://www.wbv.de/fileadmin/importiert/wbv/PDF_Website/Qualitaetsstandards_wbvOpenAccess.pdf

Großer Dank gebührt den Förderern der *wbv OpenLibrary 2025* im Fachbereich *Lehramt/Schulpädagogik*:

Fachhochschule **Münster** | Georg-August-Universität **Göttingen** | Goethe-Universität **Frankfurt am Main** | Humboldt-Universität zu **Berlin** | Justus-Liebig-Universität **Gießen** | **Karlsruhe** Institute of Technology (KIT) | Landesbibliothek **Oldenburg** | Pädagogische Hochschule **Freiburg** | TIB **Hannover** | Universität **Rostock** | Universitäts- und Landesbibliothek **Münster** | Universitäts- und Stadtbibliothek **Köln** | Universitätsbibliothek **Kassel** | Universitätsbibliothek **Kiel** | Zentral- und Hochschulbibliothek (ZHB), **Luzern** | Zentralbibliothek **Zürich**

Inhalt

Danksagung	11
Anmerkung zur geschlechtergerechten Sprache	13
1 Einleitung	15
2 Theoretischer Hintergrund	19
2.1 Deutsche Orthographie: Prinzipien und Erwerb	19
2.1.1 Die orthographischen Prinzipien als Grundlage der Schreibstrategiebildung	19
2.1.2 Orthographieerwerbsmodelle	26
2.1.3 Messung von Orthographiekompetenz	30
2.2 Fremdspracherwerbshypothesen und Orthographieerwerb	36
2.2.1 Von der Kontrastivhypothese bis zur Multi-Kompetenz- Hypothese und Theorie komplexer dynamischer Systeme (CDST) .	37
2.2.2 Begründung der theoretischen Rahmung	44
2.3 Das sprachliche Repertoire chinesischer Deutschlernender	47
2.3.1 Das multiple Sprachwissen chinesischer Deutschlernender	47
2.3.2 Kontrastive Aspekte als Einflussfaktor dynamischer Interaktion ...	51
2.3.3 Potenzielle Transferphänomene zwischen dem Chinesischen und Deutschen auf orthographischer Ebene	61
2.4 Forschungsstand zur Rechtschreibfehleranalyse im DaZ/DaF-Kontext ...	63
3 Methodik	73
3.1 Lernerkorpora in der Spracherwerbsforschung	73
3.1.1 Korpuslinguistische Zugänge	73
3.1.2 Bestehende fehlerannotierte DaF-Lernerkorpora	74
3.2 Erstellung des Lernerkorpus DeChiLKo	77
3.2.1 Datenerhebung	78
3.2.2 Datenaufbereitung	83
3.2.3 Annotation	85
3.3 Korpusabfragen in ANNIS	89
3.4 Methode zur Analyse der Korpusdaten	92
4 Orthographische Fehler der chinesischen Deutschlernenden	95
4.1 Überblick	95
4.2 Der phonographische Bereich	98
4.2.1 Fehler der Buchstabenform [Graph_BF]	99
4.2.2 Fehler bei der Phonem-Graphem-Korrespondenz [Graph_PGK] ...	102

4.2.3	Fehler bei der Markierung von Vokalquantität [Graph_VQM]	108
4.2.4	Fehler bei der s-Schreibung [Graph_s-Schreibung]	113
4.2.5	Fehler bei der Schreibung von „speziellen Graphemen“ [Graph_SG]	120
4.3	Der silbische Bereich	125
4.3.1	Fehler am Anfangsrand [Silben_AR]	126
4.3.2	Fehler im Silbenkern [Silben_SK]	130
4.3.3	Fehler am Endrand [Silben_ER]	136
4.4	Der morphologische Bereich	140
4.4.1	Fehler bei der morphologischen Stammkonstanz [Morph_Konstanz]	143
4.4.2	Fehler bei der Wortbildung: [Morph_Kompo] [Morph_Affix]	147
4.5	Der syntaktische Bereich	154
4.5.1	Fehler bei der Groß- und Kleinschreibung [Syn_GKS]	155
4.5.2	Fehler bei der Getrennt- und Zusammenschreibung [Syn_GZS] . . .	160
4.5.3	Verwechslung zwischen man und Mann [Syn_man-Mann] sowie dass und das [Syn_dass-das]	164
4.6	Der sonstige Bereich	167
4.6.1	Fehler bei der Fremdwortschreibung [Son_FW]	168
4.6.2	Grammatisch abzuleitende Fehler [Son_Gra]	170
4.6.3	Sonstige Fehler [SON]	174
5	Relation der Orthographieleistung zu verschiedenen Faktoren	177
5.1	Faktoren im Prüfungskorpus	177
5.1.1	Statistische Analyse mittels Hurdle-Modells	178
5.1.2	Der Faktor <i>Diktattext</i>	180
5.1.3	Der Faktor <i>Region</i>	183
5.1.4	Der Faktor Prüfungsstatus	185
5.1.5	Der Faktor Hochschulkategorie	188
5.2	Einflussfaktoren für die Orthographiekompetenz im Erwerbskorpus	191
5.2.1	Statistische Analyse mittels Hurdle-Modells	191
5.2.2	Der Faktor Lerndauer	192
6	Ergebnisse und Diskussion	201
6.1	Zusammenfassung und Diskussion der Ergebnisse	201
6.1.1	Spezifische orthographische Herausforderungen und Transferphänomene	201
6.1.2	Die Dynamik des Orthographieerwerbs	208
6.1.3	Das Zusammenspiel der Einflussfaktoren	210
6.2	Didaktische Implikationen	214
6.3	Methodische Limitationen der Studie	219
7	Schluss	221

Literaturverzeichnis	223
8 Anhang	243
8.1 Informationen zu den ausgewählten Hochschulen des DeChiLKo- Prüfungskorpus	243
8.2 Transkripte der Diktate im DeChiLKo-Prüfungskorpus	244
8.2.1 Transkripte der Diktat-Audiodatei von der PGG-2017	244
8.2.2 Transkripte der Diktat-Audiodatei von der PGG-2019	245
8.3 Verfügbare Metadaten im DeChiLKo-Korpus	245
8.4 Studienverlaufsplan der Germanistikausbildung an der Anhui-Universität	248

Danksagung

An dieser Stelle möchte ich all jenen danken, die zum Gelingen dieser Dissertation beigetragen und mich auf dem langen Weg der Promotion begleitet und unterstützt haben.

Mein tiefster Dank gilt meiner Betreuerin und meinen Betreuern, Frau Prof. Dr. Andrea Bogner, Herrn Prof. Dr. Marco Coniglio und Herrn Dr. Hagen Hirschmann. Ihre wissenschaftliche Begleitung, die anregenden Diskussionen und ihr stetes Vertrauen in meine Arbeit waren für mich von unschätzbarem Wert. Für die kritischen und zugleich stets konstruktiven Ratschläge, die entscheidend zur Weiterentwicklung dieser Arbeit beigetragen haben, bin ich ihnen außerordentlich dankbar.

Ein besonderer Dank gilt zudem der anonymen Gutachterin bzw. dem anonymen Gutachter der Reihe „Thema Sprache“ für die sorgfältige Lektüre des Manuskripts und die überaus konstruktiven Hinweise, die maßgeblich zur qualitativen Verbesserung dieser Publikation beigetragen haben.

Ein besonderer Dank gebührt Frau Prof. Dr. Kong Deming von der Universität Nanjing, die mir freundlicherweise die Prüfungsdaten der PGG zur Verfügung gestellt und damit eine wesentliche Grundlage für mein Prüfungskorpus geschaffen hat. Ebenso herzlich möchte ich mich bei Frau Wang Huiping von der Universität Anhui für ihre großartige Unterstützung und Kooperation bei der Sammlung der Daten für das Erwerbskorpus bedanken. Ohne ihre Hilfe wäre die empirische Basis dieser Arbeit nicht denkbar gewesen.

Für die wertvolle technische Unterstützung bei der Erstellung und Analyse des Korpus möchte ich mich herzlich bei Dr. Andreas Nolda und Tang Pei bedanken. Ihre Expertise war für die Umsetzung des korpuslinguistischen Teils dieser Untersuchung unerlässlich.

Ich danke dem China Scholarship Council (CSC) für die Gewährung des Promotionsstipendiums, das mir den Forschungsaufenthalt in Deutschland und die Konzentration auf mein Dissertationsprojekt ermöglicht hat.

Ein herzlicher Dank geht auch an meine Freund:innen und Kollegen in Göttingen, insbesondere an Leo, Hua, Vera und Minyue. Der wissenschaftliche Austausch, die gegenseitige Motivation und vor allem eure Freundschaft waren eine große Stütze in dieser Zeit.

Meinen Eltern danke ich für ihre unendliche Liebe, ihren unerschütterlichen Glauben an mich und ihre bedingungslose Unterstützung während meines gesamten Studiums und der Promotionszeit.

Zuletzt möchte ich meiner „Link-Familie“, Christa und Lutz Morgenstern, von ganzem Herzen danken. Sie haben mir in Göttingen ein Gefühl von Zuhause gegeben. Dafür bin ich zutiefst dankbar.

Anmerkung zur geschlechtergerechten Sprache

In der vorliegenden Dissertation wird durchgängig eine gendergerechte Sprache verwendet. Ziel ist es, alle Geschlechter und Geschlechtsidentitäten sprachlich sichtbar zu machen, anzuerkennen und respektvoll einzubeziehen.

Die Debatte um geschlechtergerechte Formulierungen im deutschen Sprachraum ist vielfältig und befindet sich in einem stetigen Wandel. Es existiert bislang kein alleiniger, universell anerkannter Standard. Um dieser sprachlichen und gesellschaftlichen Dynamik Rechnung zu tragen, werden in dieser Arbeit bewusst verschiedene Formen der gendergerechten Sprache abwechselnd genutzt. Dies geschieht als Ausdruck der Auseinandersetzung mit den unterschiedlichen Möglichkeiten und deren jeweiliger Eignung im Kontext.

1 Einleitung

Als ein wichtiger Bestandteil der deutschen Sprache ist das Thema Orthographie im Bereich Erstspracherwerb Deutsch (DaE) schon seit Jahrzehnten ein intensiv diskutierter Gegenstand. Für den schulischen Kontext existieren zahlreiche Lern- und Übungsmaterialien sowie Handreichungen zur Orthographievermittlung. Zudem stehen verschiedene etablierte Testverfahren wie die Hamburger Schreibprobe (HSP) (May, 2013), die Aachener förderdiagnostische Rechtschreibfehler-Analyse (AFRA) (Herné & Naumann, 2016), und die Oldenburger Fehleranalyse (OLFA) (Thomé & Thomé, 2020) zur Verfügung. Entsprechend wird das Thema Orthographie in der Erwerbsforschung umfassend untersucht (vgl. u. a. Blatt, 2010; Bredel & Müller, 2010; Bredel & Röber, 2011; Scheerer-Neumann, 1987; Thomé, 1999).

Auch im Bereich Deutsch als Zweitsprache (DaZ) haben die Themen Alphabetisierung sowie Schrift- und Orthographieerwerb in den letzten Jahren erheblich an Bedeutung gewonnen. Während die Alphabetisierung den grundlegenden Zugang zum Schriftsystem ermöglicht, stellt der Orthographieerwerb den notwendigen Entwicklungsprozess dar, in dem die spezifischen Regularitäten der deutschen Schriftsprache internalisiert werden. Diese orthographische Sicherheit bildet wiederum ein wesentliches Fundament für den Erwerb bildungssprachlicher Kompetenzen, das eine normgerechte Verschriftlichung im schulischen Kontext als Basisanforderung für eine erfolgreiche schriftliche Kommunikation gilt. Die zunehmende Heterogenität sprachlicher Lernvoraussetzungen macht dabei eine differenzierte Auseinandersetzung mit diesen spezifischen Anforderungen des Zweitspracherwerbs für eine wachsende Zahl von Lernenden notwendig. Dieser Bedarf an einer fundierten Vermittlung spiegelt sich sowohl in einer Vielzahl wissenschaftlicher Studien (vgl. u. a. Becker, 2013; Corvacho Del Toro, 2004; Geist & Hesselbarth, 2017; Grieshaber & Kalkavan-Aydın, 2012; Kalkavan-Aydın, 2012; Röber, 2012) als auch in der Entwicklung spezifischer Lehr- und Lernmaterialien (vgl. u. a. Halman, 2018; Nevyjel, 2020; Schneider et al., 2013; Wachendorf, 2022) wider.

Trotz der wachsenden Bedeutung im DaZ-Kontext wird der systematische Orthographieerwerb im Bereich Deutsch als Fremdsprache (DaF) im Vergleich zur DaE- und DaZ-Forschung noch immer randständig behandelt (vgl. Földes, 2000; Hans-Bianchi & Katelhön, 2010). Selbst in zentralen Referenzwerken der Disziplin, wie dem *Handbuch Deutsch als Fremd- und Zweitsprache* (Krumm et al., 2010), nimmt die Orthographiedidaktik im Vergleich zu anderen sprachlichen Teilkompetenzen einen erkennbar niedrigen Stellenwert ein: Während Kernbereichen wie Grammatik und Wortschatz jeweils mehrere differenzierte Fachbeiträge zugeteilt sind, die Erwerbsprozesse, Modelle und Vermittlungsmethoden detailliert beleuchten, wird die Orthographie lediglich in einem einzigen Artikel (Liang & Nerlich, 2009) abgehandelt. Diese Forschungslücke korrespondiert mit der Gestaltung aktueller DaF-Lehrwerke, in denen orthographische Regularitäten nur selten als systematischer Lerngegenstand integriert sind (vgl. Hans-Bianchi & Katelhön, 2010). Ein Beispiel hierfür ist das Lehrwerk *Studienweg Deutsch* (Liang & Ner-

lich, 2009), das als Standardlehrwerk in den meisten Germanistik-Studiengängen an chinesischen Hochschulen verwendet wird. Dort wird auf die deutsche Orthographie nur an zwei Stellen eingegangen: in einer kurzen Erklärung zum Dehnungs-*<h>* im Phonetikvorkurs und in einer Lektüre zur Rechtschreibreform von 2006. Die dortige Erklärung des Dehnungs-*<h>* ist zudem irreführend, da das Phänomen wie folgt beschrieben wird (Liang & Nerlich, 2009, S. 21):

Das Dehnungs-h wird nicht gesprochen. 元音后的 h 为长音标记，不发音，如: [Das h nach einem Vokal dient als Dehnungszeichen und wird nicht ausgesprochen, wie z. B.:] gehen, sehen, ziehen, Uhr, Sohn, Kahn, Fahne, ruhen, mahlen, ihnen.

Die angeführten Beispiele vermischen dabei verschiedene Phänomene: Während Wörter wie *Uhr* oder *Sohn* tatsächlich das Dehnungs-*<h>* aufweisen, illustrieren Beispiele wie *gehen*, *sehen*, *ziehen* und *ruhen* jedoch das silbeninitiale *<h>*.

Diese Vernachlässigung hat mehrere Gründe. In der DaF-Lehrpraxis wird die Rechtschreibung von Lernenden oft als nachrangig empfunden, da der Fokus der erwachsenen Ausbildung stärker auf inhaltlichen Aspekten liegt (vgl. Földes, 2000). Diese fachdidaktische Schwerpunktsetzung beeinflusst konsequenterweise auch die Konzeption von Lehrwerken: Da die Rechtschreibkompetenz im Fremdsprachenerwerb häufig als bloßes „Nebenprodukt“ des allgemeinen Spracherwerbs angesehen wurde (vgl. Hans-Bianchi, 2007; Hans-Bianchi & Katelhön, 2010), fehlt aufseiten der Verfassenden und Verlage oft die Notwendigkeit, systematische orthographiedidaktische Konzepte zu integrieren. Verstärkt wird diese Lücke dadurch, dass es in der akademischen Ausbildung von Lehramtsstudierenden und Lehrkräften mitunter an einer tiefgreifenden Auseinandersetzung mit der deutschen Schriftstruktur fehlt (vgl. N. Hofmann, 2008; Jagemann, 2015). Dies führt zu einem Kreislauf, in dem mangelndes fachliches Hintergrundwissen die didaktische Vermittlung erschwert und gleichzeitig die Nachfrage nach fundierten Lehrmaterialien in der Praxis gering bleibt. Folglich existieren bis heute nur wenige Forschungsarbeiten, die sich explizit mit dem Orthographieerwerb von DaF-Lernenden befassen.

Obwohl der Orthographie in Forschung und Unterrichtspraxis häufig ein nachgeordneter Stellenwert eingeräumt wird, stellt ihr Erwerb für Deutschlernende eine erhebliche Herausforderung dar. Die wenigen existierenden Studien zeigen, dass DaF-Lernende mit vielfältigen Hintergründen vor spezifischen orthographischen Herausforderungen stehen. Dies belegen Untersuchungen zu unterschiedlichen Erstsprachen, wie etwa dem Italienischen, dem Niederländischen, dem Slowenischen sowie punktuell auch dem Chinesischen (vgl. Chi, 2016; Gabrič, 2019; Hans-Bianchi, 2007; Hans-Bianchi & Katelhön, 2010; Lavalaye, 2015; Tang, 2003). Dabei fokussiert sich die Forschung bisher vorwiegend auf Lernende mit europäischen Erstsprachen wie Türkisch, Spanisch, Russisch oder Italienisch (vgl. u. a. Becker, 2013; 2018; Corvacho Del Toro, 2004; Hans-Bianchi, 2007; Hans-Bianchi & Katelhön, 2010; H.-G. Müller & Schroeder, 2022; Nimz, 2022; Thomé, 1987). Obwohl Lernende mit Chinesisch als L1 eine der wichtigsten Gruppen im DaF-Bereich repräsentieren (Auswärtiges Amt, 2020, S. 30), gibt es auffallend wenige Untersuchungen zu ihrem Schriftspracherwerb, insbesondere im Bereich

der Orthographie. Es stellt sich daher die dringende Frage nach der orthographischen Kompetenz und dem spezifischen Entwicklungsprozess chinesischer Deutschlerner.

Die vorliegende Arbeit zielt darauf ab, diese Forschungslücke zu schließen und die orthographische Leistung chinesischer DaF-Lernender auf Basis eines selbst erstellten **deutsch-chinesischen Lernerkorpus** (DeChiLKo) systematisch zu untersuchen. Der theoretische Rahmen der Arbeit basiert auf einer Integration der Multi-Kompetenz-Hypothese (vgl. Cook, 1992, 2012, 2016) und der Complex Dynamic Systems Theory (CDST) (vgl. de Bot, 2017; Larsen-Freeman, 1997; Verspoor, Lowie & Wieling, 2021), um sowohl die strukturelle Vernetzung des mehrsprachigen Wissens als auch die dynamischen, nicht linearen Prozesse des Erwerbs adäquat zu modellieren. Auf dieser Grundlage sollen auf Basis einer großen Datenmenge (329 Diktattexte) empirisch zuverlässige Erkenntnisse über die orthographische Kompetenz dieser spezifischen Lernergruppe gewonnen werden. Die detaillierte Multi-Ebenen-Annotation von orthographischen Abweichungen ermöglicht es, die komplexen Zusammenhänge zwischen Graphematik, Silbenstruktur, Morphologie und Syntax zu analysieren. Dabei stehen folgende Forschungsfragen im Zentrum: Welche spezifischen Herausforderungen und typischen Fehlermuster zeigen sich bei chinesischen DaF-Lernenden, und wie sind diese als Zusammenwirken innerhalb ihres multi-kompetenten Repertoires (L1 Chinesisch, L2 Englisch, L3 Deutsch) zu interpretieren? Welchen Verlauf nimmt die Entwicklung der Orthographiekompetenz in der Anfangsphase, und welche Dynamiken lassen sich dabei identifizieren? Welche externen und internen Faktoren (z. B. Hochschultyp, Region, Lerndauer) beeinflussen die Rechtschreibleistung, und wie ist ihr Zusammenspiel zu charakterisieren? Welche didaktischen Implikationen lassen sich aus der Fehleranalyse für die Förderung der orthographischen Kompetenz im DaF-Unterricht ableiten?

Die vorliegende Arbeit gliedert sich in sechs Hauptkapitel. Kapitel 2 legt den theoretischen Hintergrund dar. Zunächst werden die Prinzipien der deutschen Orthographie und zentrale Modelle des Orthographieerwerbs vorgestellt. Anschließend werden relevante Fremdspracherwerbshypothesen – von der Kontrastivhypothese bis zur Multi-Kompetenz-Hypothese und der CDST – erörtert und der theoretische Bezugsrahmen dieser Arbeit begründet. Darauf aufbauend werden das sprachliche Repertoire chinesischer Deutschlerner kontrastiv analysiert und der Forschungsstand zum Orthographieerwerb im DaF/DaZ-Kontext aufgearbeitet. Kapitel 3 beschreibt detailliert die Methodik, insbesondere die Erstellung, Aufbereitung und Annotation des DeChiLKo-Lernerkorpus sowie die quantitativen und qualitativen Analysemethoden. Kapitel 4 präsentiert die Ergebnisse der deskriptiven und qualitativen Fehleranalyse in den verschiedenen orthographischen Bereichen (phonographisch, silbisch, morphologisch, syntaktisch und sonstige). Kapitel 5 analysiert inferenzstatistisch die Relation der Orthographieleistung chinesischer Lernender zu verschiedenen Einflussfaktoren in den beiden Subkorpora. Abschließend werden in Kapitel 6 die zentralen Ergebnisse zusammengefasst, im Lichte der Forschungsfragen und des theoretischen Rahmens diskutiert, didaktische Implikationen abgeleitet, Limitationen der Studie reflektiert und ein Ausblick auf zukünftige Forschung gegeben.

2 Theoretischer Hintergrund

2.1 Deutsche Orthographie: Prinzipien und Erwerb

2.1.1 Die orthographischen Prinzipien als Grundlage der Schreibstrategiebildung

Das deutsche Schriftsystem bzw. die Wortschreibung kann auf Basis weniger Prinzipien modelliert werden (vgl. Drosdowski & Eisenberg, 1995; Eisenberg, 2013, 2017; Fuhrhop, 2021). Auf der Wortebene sind vor allem das phonographische, silbische sowie morphologische Prinzip von Bedeutung. Über das einzelne Wort hinaus existieren in Bezug auf die syntaktische Ebene die wortübergreifenden Prinzipien, die syntaxbezogene Rechtschreibung wie Groß- und Kleinschreibung, Getrennt- und Zusammenschreibung sowie die Interpunktionssetzung regeln (vgl. Drosdowski & Eisenberg, 1995; Fuhrhop, 2021). Auf Grundlage einfacher orthographischer Regularitäten, die durch wenige Prinzipien abgedeckt werden, können die Lernenden im Prinzip den weitaus überwiegenden Teil der Wörter der Zielsprache richtig schreiben (vgl. Eisenberg, 1995, S. 182–183, 2013, S. 288).

Im Folgenden werden vier orthographische Prinzipien vorgestellt. Diese sind grundlegend, um die orthographische Kompetenz von DaF-Lernenden zu messen. Hierzu müssen Fehler in der Lernersprache nicht nur erkannt, sondern auch klassifiziert und gedeutet werden (vgl. Eisenberg, 1995, S. 178).

2.1.1.1 Das phonographische Prinzip

Das phonographische Prinzip gilt als das grundlegende Prinzip der Alphabetschrift. Ein großer Teil der Wörter aus dem deutschen Kernwortschatz ist rein phonographisch geschrieben (vgl. Eisenberg, 2013, S. 289), d. h. solche Wortschreibungen beruhen ausschließlich auf der Beziehung zwischen Phonemen und Graphemen, die über Phonem-Graphem-Korrespondenzregeln (PGK-Regeln) festgelegt sind (vgl. Eisenberg, 2017, S. 36). Dementsprechend stellt die PGK-Kongruenz dar, „welches Segment des Geschriebenen einem bestimmten Phonem im Normalfall entspricht“ (Eisenberg, 2016, S. 68). Die folgende Abbildung beschreibt die grundlegenden PGK-Regeln für den deutschen Kernwortschatz (vgl. Eisenberg, 2013, S. 291–292, 2016, S. 69–70).

Tabelle 1: Grundlegende Phonem-Graphem-Korrespondenzregeln (Quelle: Eisenberg, 2016, S. 68–69)

Gespannte Vokale		Ungespannte Vokale	
/i:/ → <ie>	/kil/- <Kiel>	/ɪ/ → <i>	/mɪlç/ → <Milch>
/y/ → <ü>	/vyst/ -<wüst>	/ʏ/ → <ü>	/hʏpʃ/ → <hübsch>
/e/ → <e>	/vern/ → <wern>	/ɛ/ → <e>	/wɛlt/ → <Welt>

(Fortsetzung Tabelle 1)

Gespannte Vokale		Ungespannte Vokale	
/ø/ → <ö>	/ʃøn/ → <schön>	/œ/ → <ö>	/kœln/ → <Köln>
/ɛ:/ → <ä>	/bɛ:R/ → <Bär>		
/a:/ → <a>	/tʀa:n/ → <Tran>	/a/ → <a>	/kalt/ → <kalt>
/o/ → <o>	/ton/ → <Ton>	/ɔ/ → <o>	/fʀɔst/ → <Frost>
/u/ → <u>	/mut/ → <Mut>	/ʊ/ → <u>	/gʊʀt/ → <Gurt>
Reduktionsvokal			
/ə/ → <e>	/kɪʀçə/ → <Kirche>		
Diphthonge			
/aɪ/ → <ei>	/baɪn/ → <Bein>		
/aʊ/ → <au>	/tsaʊn/ → <Zaun>		
/ɔʏ/ → <eu>	/hɔʏ/ → <Heu>		
Konsonanten			
/p/ → <p>	/pɔst/ → <Post>	/ç/ → <ch>	/mɪlç/ → <Milch>
/t/ → <t>	/ton/ → <Ton>	/x/ → <ch>	/baχ/ → <Bach>
/k/ → <k>	/kalt/ → <kalt>	/v/ → <w>	/vɛʀk/ → <Werk>
/b/ → 	/bʊnt/ → <bunt>	/j/ → <j>	/jʊŋ/ → <jung>
/d/ → <d>	/dʊʀst/ → <Durst>	/h/ → <h>	/haʃt/ → <hart>
/g/ → <g>	/gʊnst/ → <Gunst>	/m/ → <m>	/mɪlç/ → <Milch>
/f/ → <f>	/fʀɔʃ/ → <Frosch>	/n/ → <n>	/napf/ → <Napf>
/z/ → <s>	/zamt/ → <Samt>	/ŋ/ → <ng>	/jʊŋ/ → <jung>
/s/ → <ß>	/ʀus/ → <Ruß>	/l/ → <l>	/lɪçt/ → <Licht>
/ʃ/ → <sch>	/ʃʀot/ → <Schrot>	/ʀ/ → <r>	/ʀɛçt/ → <Recht>
Phonemfolgen			
/kv/ → <qu>	/kvɛlə/ → <quelle>		
/ks/ → <x>	/bɔksən/ → <boxen>		
Affrikate			
/ts/ → <z>	/tsaɪt/ → <Zeit>		

Aus der Übersicht wird deutlich, dass es keine 1:1-Zuordnung zwischen Graphemen und Phonemen gibt. Diese Besonderheit wird insbesondere für jene PGK-Regeln wichtig, die Gegenstand der vorliegenden Untersuchung sind. Besonders auffällig ist beispielsweise die graphische Verschriftlichung der Phoneme /e/, /ɛ/ und /ə/, die alle durch das Graphem <e> realisiert werden.

Für die vorliegende Arbeit sind neben dem Reduktionsvokal /ə/ auch das sogenannte a-Schwa /ɐ/ von Relevanz, das beispielsweise im Auslaut von Wörtern auftaucht und graphisch als <er> realisiert wird (vgl. Fuhrhop & Peters, 2013, S. 58–59).

Darüber hinaus wird in dieser Arbeit die s-Schreibung thematisiert, da diese eine Besonderheit des phonographischen Systems widerspiegelt (vgl. Fuhrhop, 2021, S. 10). Die Stimmhaftigkeit der im Deutschen durch <s> repräsentierten s-Laute kann als positionsabhängig beschrieben werden (vgl. Eisenberg, 2013, S. 84–85). Im Kernwortschatz des Deutschen ist <s> im Anlaut ausschließlich stimmhaft, wie im *See* (/ze:/), während es im Auslaut nur als stimmlose Variante [s] vorkommt, wie in *Autos* (/autos/). Phonematisch sind die zwischensilbischen s-Laute unterschieden, wie etwa in *reisen* – *reißen*, *Muse* – *Muße* usw. Diese Komplexität spiegelt auch die Vielschichtigkeit des deutschen phonologischen Systems wider (vgl. Fuhrhop, 2021, S. 10) und kann somit eine Herausforderung für Deutschlernende darstellen.

Während einige Graphem-Phonem-Korrespondenzen im Deutschen stabil sind, zeigt bereits die s-Schreibung, dass eine rein klangliche Ableitung oft nicht ausreicht. Wie oben gezeigt, ist die Realisierung der s-Laute eng an deren Silbenposition geknüpft. Damit markiert die s-Schreibung bereits den Übergang zum silbischen Prinzip.

2.1.1.2 Das silbische Prinzip

Das silbische Prinzip beschreibt die Beziehung zwischen den Silbenstrukturen und der Schreibung (vgl. Bredel, Fuhrhop & Noack, 2011, S. 49–51). Wie das Beispiel der s-Schreibung bereits verdeutlicht, lässt sich die korrekte Schreibung vieler Wörter nicht allein anhand der GPK-Regeln ableiten; vielmehr müssen dabei silbische Informationen berücksichtigt werden (vgl. Eisenberg, 2016, S. 71). Das silbische Prinzip regelt demnach u. a. die Vokalqualität (Gespanntheit und Ungespanntheit von Vokalen) und auch Vokalquantität (Länge und Kürze von Vokalen). Ebenfalls folgt die Struktur von Anfangs- und Endrändern einer Silbe Regeln (vgl. A. Müller, 2010, S. 42). Analog zur Silbe in der Phonologie wird auch in der Graphematik von einer Silbeneinheit ausgegangen, jedoch weist eine graphematische Silbe (Schreibsilbe) im Vergleich zur Sprechsilbe eine höhere Formkonstanz auf (vgl. Fuhrhop & Peters, 2013, S. 216; Fuhrhop, 2021, S. 14).

Bei dem Aufbau von Silben kann von folgenden Regelmäßigkeiten ausgegangen werden: Anfangsrand (AR), Kern (K) und Endrand (ER) bilden eine Silbe (σ), Kern und der ER bilden den Reim der Silbe (vgl. Eisenberg, 2013, S. 296–297). Anfangs- und Endränder bestehen aus konsonantischen Materialien, die meistens fakultativ sind. Der Kern ist obligatorisch und besteht aus mindestens einem Vokalbuchstaben (vgl. Fuhrhop & Peters, 2013, S. 217). Der Typ des Anschlusses zwischen dem vokalischen Element und dem nachfolgenden konsonantischen Material spielt eine entscheidende Rolle (vgl.

Thelen, 2010, S. 36). Daher werden die Silbenkerne in dieser Arbeit zwischen Kernen mit festem Anschluss und solchen mit losem Anschluss unterschieden. In Anlehnung an die Klassifikation der Silbenkerne von Thelen (vgl. ebd.) werden die Abweichungen bei den Silbenkernen mit losem Anschluss in drei Unterkategorien eingeteilt: Monophthonge, Diphthonge und zentralisierende Diphthonge (öffnende Diphthonge), nämlich Kombinationen aus einem Vokal und vokalisiertem <r>, wie z. B. /i:ɐ/ in *wir*. Obwohl eine phonetische und damit schriftnähere Repräsentation /ɐ/-Phoneme grundsätzlich als Konsonanten analysieren und dem Endrand zuordnen könnte, würde eine stärker phonetisch orientierte Repräsentation die Vokalisierung mitberücksichtigen und daher das vokalisierte /ɐ/ mit zum Kern rechnen. Für die Erklärung orthographischer Leistungen ist ein Rückgriff auf phonetische Realisierungen häufig von Vorteil anstatt auf eher abstrakte phonologische Repräsentationen. Allerdings ist dann zu berücksichtigen, in welcher Art und Weise die Schreibenden die angenommene phonetische Realisierung tatsächlich wahrnehmen.

Zunächst geht dieser Abschnitt genauer auf die für die vorliegende Arbeit relevanten silbischen Schreibungen ein: die Silbengelenkschreibung und die Schreibung von silbeninitialem <h>.

Silbengelenke treten auf, wenn in der phonologischen Struktur eines Wortes ein betonter ungespannter und ein unbetonter Vokal durch einen Konsonanten getrennt werden. In der gesprochenen Sprache, beispielsweise bei /kɔmən/, gehört der Konsonant /m/ gleichzeitig zur ersten und zweiten Silbe, sodass nur ein /m/ hörbar ist. Schriftlich wird dieses Silbengelenk durch eine Verdopplung des Konsonanten als <mm> markiert, wodurch sichergestellt wird, dass beide Silben diesen Konsonanten enthalten. Die Silbengrenze liegt zwischen den beiden Buchstaben, was auch in der Worttrennung (*kom-men*) erhalten bleibt. Obwohl Silbengelenke nur nach kurzen Vokalen auftreten, bedeutet dies jedoch nicht, dass ein Kurzvokal die Ursache für die Verdopplung der Konsonanten ist; z. B. sind <man> und <mann> (hier unter Absehung der syntaktisch bedingten Großschreibung) phonologisch identisch. Das Substantiv wird mit Konsonantenverdopplung geschrieben, weil <nn> in der Pluralform <männer> ein Silbengelenk ist; beim Pronomen <man> gibt es keine zweisilbigen Flexionsformen, daher besteht kein Grund für eine Geminat (vgl. Fuhrhop, 2021, S. 19). Bei der Konsonantenverdopplung <mann> handelt es sich zwar um eine Gelenkschreibung, aber die Schreibung ist morphologisch zu begründen (siehe nächster Abschnitt). In der vorliegenden Arbeit geht es in Hinblick auf das silbische Prinzip um die Grundregularität von Silbengelenkschreibung, d. h. wenn im phonologischen Wort ein ambisilbischer Konsonant auftritt, wie in <kommen>.

Silbengelenke werden nicht in jedem Fall durch eine Konsonantenverdopplung wiedergegeben. Wenn eine Affrikate wie /tʃ/ ein Silbengelenk bildet, wird sie als <tʃ> verschriftlicht, wobei die Silbentrennung zwischen <t> und <ʃ> erfolgt. Der Nasallaut /ŋ/ wird hingegen mit <ng> dargestellt. Bei Mehrgraphen wie <sch> für /ʃ/ oder <ch> für /x/ erfolgt auch keine Verdopplung. Eine Besonderheit zeigt sich bei /k/, das als Silbengelenk durch die Schreibung <ck> gekennzeichnet wird. In den beiden zuletzt

genannten Fällen gehört der Mehrgraph zur zweiten Silbe, sodass die Silbentrennung unmittelbar davor erfolgt (vgl. Eisenberg, 2016, S. 77–78).

Das silbentrennende <h> dient der optischen Markierung der Silbengrenze und verdeutlicht somit die Wortstruktur (vgl. Blatt, Hein & Streubel, 2015, S. 8–9). Zwar wird es in der Standardaussprache oft nicht realisiert, es kann jedoch bei deutlicher Artikulation als Glottalfrikativ hörbar werden (vgl. Eisenberg, 2016, S. 76; Kleiner, 2011 ff., AADG: HimInlaut). In zweisilbigen Wörtern leitet es die zweite Silbe ein und wird daher als **silbeninitiales <h>** bezeichnet. Es tritt insbesondere dann auf, wenn zwei Vokale aufeinandertreffen würden (vgl. Budde, Riegler & Wiprächtiger-Geppert, 2011, S. 119) und zwischen dem betonten Vollsilbenvokal und dem unbetonten Reduktionssilbenvokal kein Konsonant hörbar ist (vgl. Bredel et al., 2011, S. 50). Wird der Silbenkern durch einen Diphthong gebildet, entfällt das <h> in der Regel (vgl. Blatt et al., 2015, S. 9). Nach dem Schreibdiphthong <ei> kann es jedoch erhalten bleiben, wie etwa im Wort *Reiher*. Da die Funktion des silbeninitialen <h> per definitionem an das Vorhandensein einer Silbengrenze gebunden ist, ist sein Vorkommen auf mehrsilbige Strukturen (bzw. auf mehrsilbige Flexionsformen) beschränkt und schließt sich in rein einsilbigen Wortformen aus (vgl. Eisenberg, 2016, S. 76).

2.1.1.3 Das morphologische Prinzip

Das morphologische Prinzip drückt das Prinzip der Morphemkonstanz aus, wonach ein Morphem über verschiedene Wortformen hinweg seine graphematische Gestalt bewahrt (vgl. Eisenberg, 2011, S. 94) (vgl. Fuhrhop, 2021, S. 25).

Morphologische Schreibung gilt insbesondere für Flexionsformen des gleichen lexikalischen Wortes, zum Beispiel wird die Pluralform zu *Apfel* als *Äpfel* geschrieben, obwohl die phonographische Form <epfel> auch möglich wäre. Die Verschriftlichung des Phonems /ɛ/ als Graphem <ä> ist lediglich morphologisch zu begründen, weil <ä> und <a> in deren Formen ähnlich sind. Eine solche morphologische Verwandtschaft lässt sich zwischen nicht umgelauteten Stammvokalen und den entsprechenden Umlauten wie <o> und <ö>, <u> und <ü> beobachten. Dies zeigt sich in der Pluralbildung (*Mutter* – *Mütter*), in der Verbflexion (*ich nahm* – *nähme*), in der Komparation (*groß* – *größer*) sowie in der Ableitung (*Hund* – *Hündchen*) (vgl. Fuhrhop, 2021, S. 25–26).

Darüber hinaus ist die übernommene silbische Schreibung von <h> in der Arbeit relevant. Entweder silbeninitiales <h> oder Dehnungs-<h> bleiben morphologisch erhalten. Hierbei muss zwischen dem silbeninitialen <h> (das die Silbengrenze markiert) und dem Dehnungs-<h> (das die Vokallänge kennzeichnet) unterschieden werden. Sowohl das silbeninitiale als auch das Dehnungs-<h> folgen dem Prinzip der Morphemkonstanz und bleiben in flektierten Wortformen erhalten. In vielen Verben dient das Dehnungs-<h> vielmehr als Lesehilfe dafür, dass der Vokal auch bei komplexen Endrändern von flektierten Formen wie <dehnst> stets gespannt gelesen wird (vgl. Fuhrhop, 2021, S. 27). Auch bei silbeninitialem <h> wie in <drehen> ist das <h> auch in Flexionsformen wie <drehst> erhalten, obwohl in dieser Form das <h> keine Segmentierung zwischen zwei Silbenkernen verdeutlicht (vgl. Eisenberg, 2013, S. 301), gleichzeitig kennzeichnet dieses <h> die Gespanntheit des Vokals (vgl. Fuhrhop, 2021, S. 27).

Eng damit verbunden ist die morphologische Doppelkonsonanz. Hier wird die Schreibung durch Konsonantenverdopplung oder spezifische Graphemkombinationen stabilisiert, um die Morphemstruktur zu bewahren (vgl. Fuhrhop, 2021, S. 28). Ein Beispiel hierfür ist die Flexionsform *rennt*: Der Doppelkonsonant <nn> wird aus der Grundform <rennen> beibehalten, obwohl er in dieser gebeugten Form nicht als Silbengelenk fungiert, da keine zwei Silben miteinander verbunden werden. Daher wird in dieser Arbeit zwischen phonologischer Silbengelenkschreibung (wie in *rennen*) und morphologischer Konsonantenverdopplung (wie in *rennt*) unterschieden.

Ein weiteres für die vorliegende Arbeit relevantes Phänomen ist die Auslautverhärtung. Im Deutschen gibt es die Regel, dass Obstruenten (Plosive und Frikative) im Silbenendrand immer stimmlos zu artikulieren sind (Fuhrhop, 2021, S. 28). Wenn ein Obstruent beispielsweise aus morphologischen Gründen am Silbenende steht, erfolgt die Entstimmung in der gesprochenen Sprache (vgl. Bulut, 2018, S. 16). In der geschriebenen Sprache hingegen bleibt das Graphem des stimmhaften Obstruenten erhalten, sofern dieser im Silbenanfang stimmhaft realisiert wird. Ein Beispiel dafür ist das Wort *Hund*, das trotz der lautlichen Realisierung /hʊnt/ in der Schreibung den stimmhaften Konsonanten <d> bewahrt, was sich im Plural *Hunde* zeigt.

Mithilfe des morphologischen Prinzips erkennen die Lernenden, dass die Schreibung konsequent den morphematischen Zusammenhang abbildet, selbst wenn sich die Lautung ändert.

2.1.1.4 Das syntaktische Prinzip

Das syntaktische Prinzip bezieht sich auf wortübergreifende Schreibkonventionen wie die Groß- und Kleinschreibung, die Getrennt- und Zusammenschreibung sowie auf die Schreibung bestimmter Funktionswörter wie der Konjunktion *dass* oder des Pronomens *man*. Dieses Prinzip umfasst auch die Zeichensetzung, welche in der vorliegenden Arbeit jedoch ausgeklammert wird. Da die Satzzeichen während des Diktats explizit vorgelesen wurden, lassen die erhobenen Daten keine bedeutenden Rückschlüsse auf die Interpunktionskompetenz der Lernenden zu.

Auf der Satzebene gilt bei der Großschreibung die Grundregel, dass das erste Wort eines selbstständigen Satzes oder Satzäquivalents großgeschrieben wird (§ 54, Geschäftsstelle des Rats für deutsche Rechtschreibung [Hrsg.], 2024, S. 81). Diese Fälle bereiten in der Großschreibung im Allgemeinen wenige Probleme, weil Sätze mithilfe der Interpunktion meistens deutlich markiert werden (vgl. Fuhrhop, 2021, S. 39).

Der schwierigste Teil bei der Groß- und Kleinschreibung ist die Substantivgroßschreibung (Fuhrhop, 2021, S. 53). Hier geht es nicht direkt um die Frage, was ein Substantiv ist und was keines, sondern darum zu verstehen, welche Substantivbestimmung der Regel zugrunde liegt. Aus einer Prototypentheorie kann man ableiten, dass ein prototypisches Substantiv eine Reihe von bestimmten Merkmalen hat (Fuhrhop, 2021, S. 41). Fuhrhop (vgl. 2021, S. 43–52) unterscheidet dabei „morphologische“ Kriterien und „syntaktische“ Kriterien. Bei der morphologischen Substantivbestimmung wird vor allem das Flexionsverhalten betrachtet: Substantive können nach Kasus und Numerus flektiert werden, aber nicht nach Genus. Anhand dieser Eigenschaften kann

man zwischen Substantiven und anderen Wortarten unterscheiden, da Verben nicht nach Kasus flektierbar sind und Adjektive, Artikel sowie Pronomen nach Genus flektiert werden. Bei der syntaktischen Substantivbestimmung prüft man, ob das jeweilige Wort im konkreten Satz wie ein Substantiv gebraucht wird. Sehr häufig sind Substantive Kerne von Nominalgruppen und an sich artikelfähig bzw. attributfähig. Mit all diesen syntaktischen Eigenschaften lassen sich prototypische Substantive bestimmen (vgl. Fuhrhop, 2021, S. 47).

Ähnlich wie bei der Groß- und Kleinschreibung geht es auch bei Getrennt- und Zusammenschreibung darum, ob man die Regularitäten, die der Schreibung zugrunde liegen, erfassen kann (vgl. Fuhrhop, 2021, S. 42). Die entscheidende Frage ist dabei, was genau als „graphematisches Wort“ gilt und somit als Einheit geschrieben wird (Fuhrhop, 2021, S. 54). Zur Klärung dieses Problems führt Fuhrhop (2010, S. 237) zwei Prinzipien ein, die den Wortbegriff definieren:

1. Wortbildungsprinzip: ‚Verbindungen‘ aus zwei oder mehr Stämmen werden zusammengeschrieben, wenn sie aufgrund einer Wortbildung miteinander verbunden sind.
2. Relationsprinzip: Einheiten, die syntaktisch analysierbar sind, das heißt insbesondere, die in syntaktischer Relation zu anderen Einheiten in einem Satz stehen, sind syntaktisch selbstständige Wörter. Sie werden getrennt geschrieben.

In Anlehnung u. a. an die Definition von Jacobs (2005) geht Fuhrhop (2007) von einem Kern- und Peripheriebereich aus (vgl. Fuhrhop, 2011, S. 115, 2021, S. 55). In Fuhrhops Ansatz wirken das Wortbildungsprinzip und das Relationsprinzip stets in einem Spannungsverhältnis. Der Kernbereich umfasst dabei diejenigen Strukturen, in denen die Kriterien beider Prinzipien eindeutig greifen und zu einer klaren Getrennt- oder Zusammenschreibung führen. Der Peripheriebereich ist hingegen dadurch gekennzeichnet, dass die Prinzipien unklare Aussagen treffen. Dies führt häufig zu systematischen Zweifelsfällen oder zulässigen Schreibvarianten, wie etwa bei *Eis laufen* bzw. *eislaufen* (vgl. Mesch, 2011, S. 270).

Innerhalb des Kernbereichs differenziert Fuhrhop (2011, S. 116–119) weiter zwischen „unproblematischen“ (wie z. B. *Mausefalle*) und „schwieriger zu entscheidenden“ Fällen (*teetrinkend* vs. *Tee trinkend*). Obwohl Letztere für Schreibende anspruchsvoller zu entscheiden sind, werden sie nicht der Peripherie zugeordnet, da sie den orthographischen Prinzipien folgend systematisch eindeutig klärbar bleiben. Aus didaktischer Perspektive kritisiert Fuhrhop jedoch die Erläuterung der Getrennt- und Zusammenschreibung anhand der typischen Zweifelsfälle, da solche Fälle erst verstanden werden können, wenn Lernende sie mit dem Kernbereich verinnerlicht haben (vgl. Fuhrhop, 2011, S. 126).

Auch die *dass/das*-Schreibungen sowie die Schreibung des Pronomens *man* sind in der vorliegenden Arbeit relevant. Obwohl die Wortpaare *das/dass* und *man/Mann* phonologisch jeweils identisch realisiert werden (/das/ bzw. /man/), unterscheiden sie sich in ihrer grammatischen Bedeutung und syntaktischen Funktion grundlegend. Die sichtbare Differenzierung in der Schreibung fungiert im Deutschen als wesentliche

visuelle Lesehilfe: Sie erleichtert den Lesenden die schnelle Worterkennung und die syntaktische Verarbeitung, da Wortarten und Satzstrukturen bereits auf der graphematischen Ebene eindeutig markiert werden (vgl. Fuhrhop & Peters, 2013, S. 242).

Die zuvor dargestellten orthographischen Prinzipien bilden die Grundlage der Rechtschreibstrategie, die Lernenden angeboten werden, damit sie die Phänomene richtig verschriften (vgl. H. Günther, 2010, S. 47–51). Der Begriff *Rechtschreibstrategie* wird in der Forschungsliteratur nicht einheitlich verwendet. In der vorliegenden Arbeit wird er im Sinne von Handlungsfolgen verstanden, mit deren Hilfe die vorgegebene orthographische Norm auf pragmatische Weise hergeleitet, erklärt und reflektiert werden kann. Davon abzugrenzen sind einerseits Prinzipien, die die Strukturen der Schriftsprache beschreiben (z. B. phonographisches, morphematisches Prinzip), und andererseits Normen (z. B. Groß- und Kleinschreibungsregeln), die amtlich festgelegt sind und nicht als Strategien im engeren Sinne gelten (vgl. Ossner, 2006). Je nach Kontext bezeichnet der Begriff in dieser Arbeit sowohl konkrete kognitive Operationen der Lernenden (z. B. Silbierung, Wortstammableitung) als auch didaktische Verfahren, die Lehrende zur Förderung des Regelwissens einsetzen. Diese Rechtschreibstrategien lassen sich wiederum in der Rechtschreibentwicklung wiedererkennen (vgl. Bulut, 2018, S. 19). Im Folgenden werden einige Erwerbsmodelle der Orthographie beschrieben. Da zum orthographischen Erwerbsprozess bei DaF-Lernenden nur wenige Studien vorliegen (siehe Kapitel 2.4), greift diese Untersuchung auf die grundlegenden Modelle des Erstspracherwerbs zurück. Die Frage, inwieweit diese Modelle auf mehrsprachige Lernkontexte übertragbar sind, wird dabei kritisch reflektiert.

2.1.2 Orthographieerwerbsmodelle

Die Forschung zum Orthographieerwerb im Erstspracherwerb Deutsch hat eine Vielzahl von Modellen entwickelt, die unterschiedliche Entwicklungsstadien beschreiben und erklären. Davon beschreiben die meisten Erwerbsmodelle den Lernprozess als eine Abfolge von distinkten Entwicklungsstufen, in denen der Orthographieerwerb schrittweise erfolgt (vgl. Bredel et al., 2011, S. 96). Dabei werden auf Grundlage empirischer Befunde Entwicklungsstufen bestimmt, durch die Kinder zum korrekten Schriftsprachgebrauch gelangen (vgl. Bulut, 2018, S. 20–21). Als Grundlage für die frühen Stufenmodelle gilt das Stufenmodell des Lesen- und Schreibenlernens von Frith (1985), das ursprünglich für das Englische entwickelt wurde.

In Anlehnung an das Modell von Marsh, Friedman und Desberg (1981) beschreibt Frith den Schriftspracherwerb als einen mehrstufigen Prozess, in dem sich die Strategien für das Lesen und Schreiben schrittweise weiterentwickeln. Dieses Modell unterscheidet drei Hauptphasen, die jeweils in zwei Stufen unterteilt sind. Diese Phasen sind die logographische, alphabetische und orthographische Phase. Die folgende Tabelle fasst die sechs Stufen des Modells zusammen.

Tabelle 2: Das Sechs-Stufenmodell von Frith (1985) (Quelle: Frith, 1985, S. 311)

Step	Reading	Writing
1a	logographic ₁	(symbolic)
1b	logographic ₂	logographic ₂
2a	logographic ₃	alphabetic ₁
2b	alphabetic ₂	alphabetic ₂
3a	orthographic ₁	alphabetic ₃
3b	orthographic ₂	orthographic ₂

Es wird im Modell eine wechselseitige Förderung zwischen Lesen und Schreiben konstruiert, wobei die Pfeile auf Übergänge zur nächsten Phase in den beiden Modalitäten hindeuten (vgl. Frith, 1985, S. 311). In der logographischen Phase erkennen Kinder Wörter anhand visueller Merkmale wie Schrifttyp oder Farbe der Schrift. Orthographische Details wie die Buchstabenreihenfolge spielen keine wesentliche Rolle. Kinder erkennen bekannte Wörter wieder, die in ihrem Alltagsleben häufig vorkommen, ohne Bezug auf phonologische Aspekte zu nehmen. Für das Schreiben spielt die logographische Strategie kaum eine Rolle: „Use of logographic strategy would mean that the word cannot be written at all“ (Frith, 1985, S. 313).

In der alphabetischen Phase beginnen Kinder, Graphem-Phonem-Korrespondenzen systematisch zu nutzen. Sie erkennen, dass Buchstaben für Laute stehen, und beginnen, Wörter Laut für Laut zu erschließen. Für das Schreiben versuchen die Kinder, für die lautlichen Einheiten entsprechende schriftliche Einheiten zu finden und zu schreiben; im idealen Fall entstehen auf diese Weise lautgetreue Schreibungen (vgl. Siekmann & Thomé, 2018, S. 117). Da jedoch die orthographischen Regeln noch nicht vollständig verinnerlicht sind, treten in dieser Phase viele lautgetreue, aber orthographisch nicht normgemäße Schreibungen auf.

In der orthographischen Phase erkennen Kinder zunehmend größere sprachliche Einheiten, wodurch das Lesen und Schreiben flüssiger und automatisierter wird (Frith, 1985, S. 309). Während in der alphabetischen Phase die Verarbeitung einzelner Buchstaben dominiert, ermöglicht die orthographische Strategie das Erkennen von immer wieder vorkommenden Buchstabenkombinationen, häufigen Wörtern oder Morphemen. Für das Deutsche ist dieser Übergang insbesondere durch die Beachtung der Stammkonstanz und silbischer Muster gekennzeichnet. Das Lesen wird dadurch schneller, da gespeicherte Wortmuster aus dem mentalen Lexikon abgerufen werden. Beim Schreiben sind Kinder nicht mehr ausschließlich auf die Laut-Buchstaben-Zuordnung angewiesen, sondern nutzen orthographische Regeln und Muster. In dieser Phase treten die Normabweichungen auf, die „unangebrachte morphemorientierte Verschriftungen oder die Verwendung komplexer Grapheme als im gegebenen Wort notwendig“ sind, wie etwa die Fehlschreibung *<wright> für *write* (Siekmann & Thomé, 2018, S. 118).

Im Gegensatz zu dem rein additiven Erwerbsprozess, der den meisten Stufenmodellen theoretisch zugrunde liegt (vgl. Bredel et al., 2011, S. 96; Siekmann & Thomé, 2018, S. 118), betrachtet Frith (1985) den Schriftspracherwerb als einen dynamischen Prozess der Verschmelzung einzelner Strategien. Dies bedeutet, dass eine neue Strategie nicht einfach die vorherige ablöst, sondern dass der Erwerb vielmehr durch einen komplexen Prozess der Reorganisation und Integration geprägt ist. Die jeweils ältere Strategie bleibt latenter Bestandteil der neuen und wird in latenter Form weitergeführt. Ähnliche Annahmen werden in anderen Entwicklungsmodellen als „hierarchische Parallelität“ (Thomé, 2006, S. 375) bzw. „zeitlich versetzte Parallelität“ (Eichler & Thomé, 1995, S. 35) bezeichnet. Dabei dominiert eine Schreibstrategie, während die vorhergehenden Entwicklungsstufen latent erhalten bleiben und bei Bedarf verwendet werden können (vgl. Thomé, 2006, S. 375).

Viele Wissenschaftler*innen entwickelten u. a. auf Grundlage der dargestellten Stufenmodelle eine Vielzahl von Schriftsprachentwicklungsmodellen zur deutschen Sprache (vgl. u. a. Brügelmann & Brinkmann, 1994; K. B. Günther, 1986; May, 1990; Scheerer-Neumann, 1987; Thomé, 2006; Valtin, 1994). Siekmann und Thomé (2018), Thomé (2006). Jaeuthe (2023) erstellen einen umfangreichen Überblick über alle vorhandenen Schriftspracherwerbsmodelle. In der Gegenüberstellung aller Modelle lassen sich viele Gemeinsamkeiten feststellen. Obwohl die Anzahl der Entwicklungsstufen von drei (u. a. Frith, 1985; May, 1990; Thomé, 2006) bis sechs (u. a. Scheerer-Neumann, 1987; Valtin, 1994) variiert und die Bezeichnungen der einzelnen Stufen unterschiedlich als „Phasen“, „Stufen“, „Strategien“, „Level“, „Regeln“, „Etappen“, „Stadien“, „Prinzipien“, „Teilfähigkeiten“ usw. dargestellt werden, besteht weitgehender Konsens über die grundlegende Entwicklungsabfolge: Diese verläuft von einer primär phonographischen Orientierung hin zur Anwendung komplexer silbischer und morphologischer Regularitäten (vgl. Bredel et al., 2011, S. 95; Bulut, 2018, S. 24; Jaeuthe, 2023, S. 829). Jaeuthe fasst die qualitativen Veränderungen in allen Entwicklungsmodellen als drei Kompetenzniveaus zusammen (Jaeuthe, 2023, S. 829):

1. Es entstehen erste Schreibprodukte, wobei die Schreibungen noch nicht der Lautung eines Wortes entsprechen. Je nach Modell können diese ersten gekritzelten Zeichen oder einzelne Buchstaben sein.
2. Es entstehen Schreibprodukte, die sich an der Lautung eines Wortes orientieren, sodass die Schreibung eines Wortes seiner Lautung entspricht.
3. Es entstehen korrekte Schreibungen, welche sich nicht ausschließlich über die Lautung des Wortes erklären lassen.

Ein zentraler Kritikpunkt an orthographischen Entwicklungsmodellen betrifft die methodische Schwierigkeit, einzelne Lernende eindeutig einer bestimmten Stufe zuzuordnen. So wird im Modell von Scheerer-Neumann versucht, Schüler*innen über die mehrheitlich genutzte Strategie einem Kompetenzniveau zuzuordnen, während in den Modellen von Günther, Brügelmann und Brinkmann, Thome und May die Dominanz einer Phase betont wird. Beide Ansätze können jedoch nicht eindeutig klären,

wie Lernende zu klassifizieren sind, die zwei oder mehr Strategien mit ähnlicher Häufigkeit verwenden.

Kritisiert wird an den Stufenmodellen zudem die Annahme, die Entwicklungsphasen seien universell und würden von allen Kindern kulturunabhängig in derselben Reihenfolge durchlaufen. Dieser Universalitätsanspruch wird jedoch von zahlreichen empirischen Studien widerlegt, die zeigen, dass Entwicklungsprozesse individuell geprägt sind (vgl. Brinkmann, 2003; Bulut, 2018; May, 1990), „dass Kinder sich – je nach Kontext – entsprechend unterschiedlichen Stufen verhalten können“ (Scheerer-Neumann, 2006, S. 514). Der Prozess des Orthographieerwerbs wird von einer Vielzahl heterogener Faktoren beeinflusst, welche die individuelle Lernentwicklung maßgeblich prägen. Diese Einflussfaktoren lassen sich primär in individuelle, familiäre sowie institutionelle (schulische) Determinanten unterteilen (vgl. Bulut, 2018, S. 27). So wurde beispielsweise der signifikante Zusammenhang zwischen der kognitiven Leistungsfähigkeit der Lernenden und ihrer Rechtschreibleistung bereits in vielen empirischen Studien belegt (vgl. u. a. Becker, 2013, S. 228; Corvacho Del Toro, 2013, S. 194).

Die für den monolingualen Erstspracherwerb konzipierten Entwicklungsmodelle sind nicht ohne Weiteres auf DaZ- und DaF-Lernende übertragbar. Zwar liegen bislang nur wenige Studien vor, die die Rechtschreibentwicklung von monolingualen und bilingualen Kindern systematisch vergleichen, doch der Einfluss von Faktoren wie dem L1-Transfer auf den orthographischen Erwerbsprozess ist belegt (eine ausführliche Darstellung des Forschungsstands bietet das Kapitel 2.4). Im Gegensatz zum Erstspracherwerb verfügen DaF-Lernende zudem bereits über eine (oder mehrere) durchlaufene Alphabetisierung(en) und somit über etablierte Schreibstrategien für ihre Erstsprache. Aus diesem Grund setzt der Lernprozess im Schriftspracherwerb bei ihnen auf einer anderen Stufe an als bei Kindern in der Erstalphabetisierung (vgl. Hans-Bianchi, 2007; Hans-Bianchi & Katelhön, 2010).

Eine der wenigen Untersuchungen über den Orthographieerwerb in der Fremdsprache Deutsch von Hans-Bianchi geht von der Beobachtung der Schreibleistungen von italienischsprachigen Studierenden aus, dass die Schreibungen sehr häufig aus dem Rahmen der üblichen Phasenmodelle abweichen. Eine Fehlschreibung wie *<var-schainlig> für *wahrscheinlich* verdeutlicht dies: Einerseits enthält sie typische „Anfängerfehler“ wie etwa <ai> statt <ei> oder das fehlende Anlaut-<h>, was eine fehlerhafte alphabetische Strategie zeigt; andererseits zeugt die korrekte Wiedergabe komplexer orthographischer Muster (Suffix <-ig>; Umlautdiphthong <äu>) bereits von einer fortgeschrittenen morphematischen Strategie (vgl. Hans-Bianchi & Katelhön, 2010, S. 37).

Ein solches Schreibprofil verdeutlicht somit die Grenzen der für den Erstspracherwerb entwickelten Stufenmodelle im DaF-Kontext: Ihre Unfähigkeit, die für mehrsprachige Lernende typische, simultane Nutzung von Strategien aus unterschiedlichen Phasen zu erklären, macht eine direkte Übertragung problematisch. In der vorliegenden Untersuchung wird anhand von Korpusdaten ein detailliertes Bild der spezifischen orthographischen Entwicklung von chinesischen DaF-Lernenden sowie des Einflusses interner und externer Faktoren gezeichnet, um die Komplexität des Erwerbsprozesses dieser Lernergruppe darzustellen.

2.1.3 Messung von Orthographiekompetenz

Fehler stellen einen wesentlichen Bestandteil des Lernprozesses dar und werden seit der kognitiven Wende in der Psychologie als produktive Indikatoren für die Lernentwicklung betrachtet (vgl. Corvacho Del Toro & Lorenz, 2021, S. 1). Daher liegt gängigen Phasenmodellen des Schriftspracherwerbs die Annahme zugrunde, dass typische Fehler einer bestimmten Entwicklungsphase zugeordnet werden können (vgl. Corvacho Del Toro, 2016, S. 399; Siekmann & Thomé, 2018, S. 147). Sehr häufig werden Rechtschreibfehler durch Diktate ermittelt, wie zahlreiche Studien zur Rechtschreibleistung von Lernenden dokumentieren. Die DESI-Studie (Deutsch Englisch Schülerleistungen International) ist eine solche umfangreiche Studie. Sie wurde in den Jahren 2003/2004 durchgeführt und erhob Daten von über 11.000 Schüler*innen der 9. Jahrgangsstufe (DESI-Konsortium, 2008). Die Daten zu unterschiedlichen Teilbereichen ihrer Deutsch- und Englischkompetenz werteten die orthographische Kompetenz mittels eines anspruchsvollen Diktats mit orthographisch schwierigen Wörtern aus (vgl. Nimz, 2022; Thomé & Gomolka, 2007).

Nicht nur in der Erstsprache Deutsch, sondern auch im Rahmen des Fremdsprachunterrichts wird das Diktat häufig als didaktisches Instrument im Rechtschreibunterricht und als Werkzeug zur Beurteilung der Lernfortschritte eingesetzt. Allerdings steht diese Methode aufgrund ihrer häufig rein quantitativen Ausrichtung seit Längerem in der Kritik. Lernende erhalten oft nur allgemeine Rückmeldungen in Form einer Gesamtpunktzahl oder reinen Fehlerzahl, ohne differenzierte Hinweise auf konkrete Fehlerquellen (vgl. Siekmann & Thomé, 2018, S. 149). Eine solche rein quantitative Bewertung kann nur unzureichende Aussagen über die tatsächliche Kompetenz der Lernenden liefern, deshalb gewinnen qualitative Fehleranalysen zunehmend an Bedeutung (vgl. Corvacho Del Toro, 2016; Corvacho Del Toro & Lorenz, 2021). Ziel der qualitativen Fehleranalyse ist, „ein möglichst genaues Bild vom Stand der Rechtschreibkenntnisse einzelner Schüler zu zeichnen, indem nicht das Wort als Ganzes, sondern jede einzelne Stelle, an der Fehler auftreten können, auf Korrektheit überprüft wird“ (Siekmann & Thomé, 2018, S. 149). Solche Analysen bieten einen tieferen Einblick in Fehlererscheinungen und deren spezifische Ursachen. Dadurch können qualitative Fehleranalysen eine Basis für konkrete Unterrichts-, Förder- und Therapieplanung verschaffen (vgl. Corvacho Del Toro, 2016, S. 398 f.; Siekmann & Thomé, 2018, S. 149).

Zur Feststellung der Rechtschreibfähigkeit existieren verschiedene Verfahren, die durch qualitative Fehleranalysen eine detaillierte Einschätzung des individuellen Entwicklungsstandes, des orthographischen Wissens sowie der angewandten Lernstrategien ermöglichen (vgl. Siekmann & Thomé, 2018, S. 149). Siekmann und Thomé (2018) unterscheiden hierbei grundsätzlich zwischen testabhängigen und testunabhängigen Diagnoseinstrumenten. Zu den bekanntesten testabhängigen Verfahren zählt beispielsweise die Hamburger Schreibprobe (HSP). Die HSP ist ein standardisiertes Verfahren, das eine quantitativ-qualitative Auswertung von Schreibungen vorsieht und somit nicht nur den aktuellen Stand des Rechtschreibkönnens abbildet, sondern auch das verwendete orthographische Strukturwissen und die individuellen Rechtschreibstrategien erfassen soll (vgl. Bredel et al., 2011, S. 163). Der Test basiert auf dem Strategiemodell von

May (1990), das sich konzeptionell und terminologisch an das Modell von Frith (1985) anlehnt. Nach diesem Modell werden im Laufe des Orthographieerwerbs die folgenden Strategien entwickelt (May, 2013, S. 21):

- die logographemische,
- alphabetische,
- orthographische,
- morphematische,
- wortübergreifende Strategie

Für die Teilnahme an HSP benötigt jede*r Schüler*in ein Testheft. Je nach Klassenstufe erhalten die Kinder unterschiedliche Testinhalte. Beispielsweise umfasst HSP für die zweite Klasse 15 Einzelwörter, welche die Schüler*innen anhand der gegebenen Illustrationen verschriftlichen sollen, sowie drei Lückensätze, die den Schüler*innen diktirt werden (vgl. May, 2013, 110 f.).

Die Auswertung der Testergebnisse erfolgt in drei Teilen: wortbezogene Auswertung; Auswertung nach Graphemtreffern und nach Rechtschreibstrategien. Durch den Vergleich der Anzahl richtig geschriebener Wörter sowie der Anzahl richtig geschriebener Grapheme (sog. Graphemtreffer) wird der individuell-orthographische Lernstand der Teilnehmenden mit der Vergleichstabelle ermittelt. Exemplarisch demonstriert May (2013, S. 18) eine idealisierte Lernabfolge am Wort Fahrrad (siehe Tabelle 3). Damit wird die Annahme dargestellt, dass der Schriftspracherwerb als ein nacheinander ablaufender Prozess verschiedener Entwicklungsstufen zu verstehen ist.

Tabelle 3: Unterschiedliche Maßstäbe für die Bewertung am Beispiel verschiedener Schreibungen des Wortes „Fahrrad“ nach May (2013) (Quelle: May, 2013, S. 18)

Schreibung	Richtig/falsch	Realisiertes Rechtschreibwissen	Graphemtreffer
FT	falsch	Alphabetisches Schreiben: verkürzte Lautfolge	1
Fart	falsch	Alphabetisches Schreiben entfaltet	2
Farat	falsch	Vollständiges alphabetisches Schreiben	3
Farad	falsch	Orthographisches Regelement bezeichnet	4
Fahrad	falsch	Orthographisches Merkelement bezeichnet	5
Fahrrad	richtig	Kompositum morphematisch durchdrungen	6

Folgt man diesem idealisierten, streng linearen Modell, dann würden Fehlschreibungen wie *<Farrat> (Kompositumelement <rr>, keine orthographische Strategie) oder *<Fahrat> (orthographisches Markelement <ah>, aber nicht regelhaftes Element <rad>) eher unwahrscheinlich erscheinen. Tatsächlich sind jedoch solche „gemischten“ Fehlschreibungen in der Lernaltersprache keine Seltenheit, sondern lassen sich empirisch regelmäßig beobachten (vgl. Siekmann & Thomé, 2018, S. 170; auch im Korpus DeChiLKO). Diese Diskrepanz zwischen dem theoretischen Modell und den empiri-

schen Befunden verdeutlicht, dass der Erwerbsprozess weitaus komplexer als eine einfache Abfolge von Stufen ist.

Die Auswertung nach Rechtschreibstrategie in der HSP erfolgt anhand vordefinierter Lupenstellen, d. h. spezifischer Positionen in Wortschreibungen, die als Indikatoren für die verwendeten Rechtschreibstrategien herangezogen werden. Um diese Zuordnung für die Schulpraxis praktikabel und einheitlich zu gestalten, unterliegt die Auswahl der Lupenstellen in der HSP strengen testinternen Richtlinien. Ein zentraler Grundsatz lautet dabei: „Eine Wortstelle wird nur einer einzigen Rechtschreibstrategie zugeordnet“ (May, 2013, S. 33). Dadurch sollen mögliche Strategie-Interferenzen bei der Auswertung vermieden werden. Zudem sieht die HSP-Logik vor, aus dem Wortmaterial gezielt nur jene Phänomene als Lupenstellen auszuwählen, die unterschiedliche Strategien in ausgewogener Form abfordern. Vor diesem methodischen Hintergrund erklären sich die spezifischen Setzungen der Lupenstellen im Testmaterial. Beispielsweise wird die Dehnungsschreibung <ah> im Wort *Fahrrad* in der HSP 2 als orthographisches Merkelement gewertet und entsprechend der orthographischen Strategie zugeordnet (siehe Tabelle 3). Dass hingegen bei Wörtern wie *Schubkarre* (HSP 3) nur das als Indikator für die morphologische Strategie erfasst wird, während die Konsonantenverdopplung <rr> als Lupenstelle unberücksichtigt bleibt, ist keine Inkonsistenz. Es handelt sich vielmehr um eine bewusste testökonomische Auswahl. Ähnliches fällt für das Wort *Schlüsselloch* auf: Obwohl die <ss>-Schreibung laut den HSP-Richtlinien ein typisches Regelement der orthographischen Strategie darstellt, fungiert sie in diesem konkreten Fall nicht als Lupenstelle (ebd., S. 162), da die Anzahl der zu überprüfenden Stellen pro Wort im Sinne der Testökonomie und Auswertbarkeit begrenzt ist.

Zu den gängigen testunabhängigen diagnostischen Instrumenten gehören die Aachener Förderdiagnostische Rechtschreibfehler-Analyse (AFRA) und die Oldenburger Fehleranalyse (OLFA). Im Vergleich zu der HSP, die nur indirekt eine Fehlertypologie liefert, werden in der AFRA und OLFA konkrete Fehlerkategorien entwickelt.

AFRA wurde von den Sprachwissenschaftlern Karl-Ludwig Herné und Carl Ludwig Neumann entwickelt. Bei dem Kernkonzept der AFRA handelt es sich um ein „geordnetes Inventar linguistisch fundierter Rechtschreibfehler-Kategorien“ (Herné & Naumann, 2016, S. 6). Jeder Fehler bzw. jede Fehlerstelle im Wort wird als Verstoß gegen eine Rechtschreibregel betrachtet. Dabei versteht sich die „AFRA als deskriptiv orientiertes Verfahren, bei dem Ursachenvermutungen zugunsten der rechtschreibsystematischen Beschreibung von Fehlern zurücktreten“ (Herné, 2014, S. 142). Die Rechtschreibfehler der Schüler*innen werden mittels AFRA nicht nur quantitativ, sondern auch qualitativ ausgewertet. Die AFRA ist vor allem mit standardisiertem Rechtschreibtest durchzuführen, kann aber auch ergänzend auf Diktate und freie Texte angewandt werden.

Als Grundlage und schematische Visualisierung der orthographisch-systematischen Konzeptionen Hernés & Naumanns fungiert das „Haus der Orthographie“ (vgl. Herné & Naumann, 2016, S. 16). Diese Struktur illustriert in simplifizierter Weise die drei Analysebereiche, die den tradierten linguistischen Beschreibungsebenen der Phonologie, Morphologie und Syntax entsprechen und dabei in entwicklungstheoretischer

Sichtweise die unterschiedlichen Stufen, aber auch Schwierigkeitsgrade des Schrift-spracherwerbs repräsentieren.

Insgesamt lassen sich 16 Fehlerkategorien in AFRA einordnen (vgl. Herné & Naumann, 2016, S. 6). Herné und Naumann unterscheiden zudem, ob eine Schreibung zur Mehrheits- oder zur Minderheitsschreibung zu zählen ist. In AFRA werden Mehrheits-schreibungen durch ein Pluszeichen (z. B. LI+ für die Mehrheitsschreibung <ie> beim lang gesprochenen /i:/ wie in *sieben*) und Minderheitsschreibungen durch ein Minuszeichen gekennzeichnet (z. B. LI- für den entsprechenden Minderheitsfall <i>, <ih> oder <ieh> beim lang gesprochenen /i:/ wie etwa in *Tiger*, *sieht* und *ihnen*) (vgl. Herné & Naumann, 2016, S. 7). Diese Fehlerkategorien bündeln sich zu einer 4-Ebenen-Analyse:

- Phonologische Ebene
- Vokalquantität
- Morphologische Ebene
- Syntaktische Ebene

Nach der Durchführung eines standardisierten Rechtschreibtests ergeben sich in den Auswertungsrastern Fehlerprofile, indem Prozentwerte aus den tatsächlichen Fehlern mit den gegebenen Basisraten in einer Kategorie errechnet werden.

Die Gesamtzahl der in einer bestimmten Kategorie möglichen Fehlererscheinungen wird als Basisrate bezeichnet, nämlich die „theoretische Fehlerverlockung“ (Herné & Naumann, 2016, S. 16). Anhand der Fehlerprofile lassen sich die Bereiche verorten, in denen Lernende noch Schwierigkeiten haben. Die Möglichkeit, mehrere Vorkommen von Fehlern bei Buchstabenform (BF), Graphem-Auswahl (GA) und Graphem-Folge (GF) in einem Wort zu signalisieren, wird jedoch von den Basisraten nicht berücksichtigt, weil deren Anzahl nicht vorhersagbar ist. Das Grundkonzept der AFRA zielt zwar darauf ab, jede mögliche Fehlerstelle zu kennzeichnen, die tatsächliche Bewertung konzentriert sich jedoch nur auf die sogenannten „Lupenstellen“ (in Anlehnung an May (1990)). Hierbei handelt es sich um spezifische Wortpositionen, die als Indikatoren für bestimmte Rechtschreibstrategien dienen. So wird beispielsweise die Fehlschreibung *<Warheit> für *Wahrheit* als Verstoß gegen Morphem-Differenzierung (MD) in AFRA markiert, obwohl die Verwechslung <a> für <ah> auch als Abweichung bei der Vokalquantitätsmarkierung zugeordnet werden kann.

Außerdem wird die praktische Durchführbarkeit der AFRA mit Diktaten oder freien Texten infrage gestellt. Obwohl die Methode dank ihrer differenzierten Fehlerkategorien eine sehr genaue Analyse ermöglicht (vgl. Siekmann & Thomé, 2018, S. 183), ist ihre Anwendung mit erheblichem Aufwand verbunden. So ist bereits die Bestimmung der „Fehlerverlockung“ sehr zeitintensiv (vgl. Herné & Naumann, 2016, S. 17). Vor allem aber setzt die Analyse hohe linguistische Kenntnisstände bei den Lehrkräften voraus (vgl. Siekmann & Thomé, 2018, S. 183). Denn um die Fehler in einem freien Text (oder Diktat) systematisch zu erfassen, müssen die Lehrkräfte nicht nur jede einzelne Abweichung erkennen, sondern diese auch präzise den jeweiligen Fehlerkategorien in AFRA zuordnen können, was ein tiefes linguistisches Verständnis erfordert.

Die OLFA wurde vom Sprachwissenschaftler Günther Thomé und der Pädagogin Dorothea Thomé zusammen entwickelt. Dieses Verfahren zur qualitativen Fehleranalyse basiert auf einer umfangreichen Datensammlung aus mehreren DFG-Forschungsprojekten (vgl. Eichler & Thomé, 1995; Thomé, 1999). Ziel dieser Analyse ist es, den individuellen Lernstand der Lernenden zu erfassen und daraus eine gezielte Förderplanung abzuleiten (vgl. Siekmann & Thomé, 2018, S. 188). Im Gegensatz zu anderen Verfahren, die auf standardisierten Diktaten oder Lückentexten beruhen, basiert OLFA 3–9 auf der Analyse frei verfasster Texte, um ein möglichst authentisches Bild der orthographischen Kompetenzen zu erhalten (vgl. Thomé & Thomé, 2020, S. 11).

Die OLFA bewertet nicht nur die Anzahl der Fehlererscheinungen, sondern insbesondere die Qualität der Rechtschreibfehler. Durch die Zuordnung der Fehler zu bestimmten Schriftspracherwerbsphasen lassen sich gezielte Aussagen über die Entwicklungsstufe eines Lernenden treffen (vgl. Thomé & Thomé, 2020, S. 8). Das zugrunde liegende Entwicklungsmodell des Rechtschreibens von OLFA orientiert sich an dem Modell von Thomé (2006), das drei Phasen umfasst: die proto-alphabetisch-phonetische, die alphabetische und die orthographische Phase. Dieses Modell unterscheidet sich von den anderen Stufenmodellen darin, dass „es nicht versucht, Stufen einer mehr oder weniger linearen Entwicklung der Rechtschreibfähigkeit zu skizzieren, sondern Kategorien nennt, in die auftretende Fehler von Lernern eingeordnet werden können“ (Siekmann & Thomé, 2018, S. 141). Außerdem basiert das Konzept von OLFA auf Ergebnissen über statische Verhältnisse zwischen Phonemen und Graphemen (vgl. Thomé, 1987, 39 f.). Dabei wird zwischen Basisgraphem und Orthographem unterschieden: Als *Basisgrapheme* werden die grundlegenden und häufigsten Repräsentationen eines Phonems bezeichnet; statistisch seltener und somit oft lernschwierigere Grapheme werden hingegen als *Orthographeme* definiert (vgl. Thomé & Thomé, 2020, S. 9).

Die Analyse der OLFA erfolgt schrittweise: Zunächst werden die Fehler identifiziert und dokumentiert, anschließend in 37 vordefinierte Fehlerkategorien eingeordnet. Dabei erfolgt die Fehlersortierung nach Rechtschreibbereichen auf der horizontalen Ebene rein deskriptiv, ohne ätiologische Interpretation (vgl. Siekmann & Thomé, 2018, S. 194). Die 37 Fehlerkategorien gliedern sich in acht übergeordnete Bereiche:

- Groß- und Kleinschreibung,
- Getrennt- und Zusammenschreibung,
- Kürze- und Längenmarkierung,
- s/ß-Schreibung,
- spezielle Graphembereiche (z. B. e/eu, ä/äu),
- Stammschreibungen bei Auslautverhärtung und -velarisierung,
- f/v/w-Schreibungen sowie
- lautliche Fehler

In der vertikalen Dimension des Auswertungsrasters werden die identifizierten Fehler den drei Schriftspracherwerbsphasen (Gruppe I, II und III) zugeordnet (vgl. Thomé & Thomé, 2020, S. 25). Durch die Addition der Fehlerwerte innerhalb jeder Gruppe ergeben sich die Teilsommen, aus denen schließlich die Gesamtfehlersumme, der indivi-

duelle Fehlerwert (Fehler pro 100 Wörter) sowie die prozentualen Anteile der drei Gruppen resultieren.

Mithilfe der OLFA-Ergebnisse werden zwei zentrale Werte bestimmt, die eine differenzierte Einordnung der Rechtschreibfähigkeiten ermöglichen: der Kompetenzwert (KW) und der Leistungswert (LW). Der Kompetenzwert (KW) ist ein qualitativer Wert, der das zugrunde liegende orthographische Strukturwissen der Lernenden abbildet, unabhängig von der reinen Fehleranzahl. Ein hoher KW zeigt an, dass die grundlegenden orthographischen Bereiche bereits gefestigt sind, auch wenn in höheren Entwicklungsstufen noch Fehler auftreten. Der Leistungswert (LW) hingegen ist ein quantitativer Normwert. Er misst die sichtbare Schreibleistung und erlaubt einen direkten Vergleich mit der durchschnittlichen Leistung der jeweiligen Klassenstufe, um einen eventuellen Förderbedarf zu bestimmen (vgl. Thomé & Thomé, S. 31–36).

Ausgehend vom Forschungsinteresse der vorliegenden Arbeit werden in der Tabelle 4 die Fehlertypologien der AFRA sowie OLFA gegenübergestellt. Beide Verfahren dienen der deskriptiven Kategorisierung von Rechtschreibfehlern, unterscheiden sich jedoch in ihrer methodischen Ausrichtung. Während die AFRA eine stärker linguistische Perspektive einnimmt, ordnet OLFA die Fehler nach kognitiven Erwerbsstufen der Schriftsprachentwicklung ein. Um die Kategorisierung beider Verfahren zu illustrieren, wird exemplarisch die hierfür konstruierte Fehlschreibung *<faratchlos> für das Wort *Fahrradschloss* analysiert. Die folgende Tabelle zeigt an, wie dieser Fehler in den jeweiligen Systemen klassifiziert wird:

	faratchlos							
	f	a	r	a	t	ch	l	o
AFRA	GK- fälschliche Kleinschreibung	LV- Fehler bei der Schreibung eines durch Dehnungs-<h> gekennzeichnet langen Vokals	MS fehlerhafte Verschriftung eines Morphemanschlusses		KA+ Nichtbeachtung der Verlängerungsregeln bei <d/> usw.	SG+ Fehler bei <ch>, <ß>, <k> usw.		KV ₁₀ + Fehler bei der Schreibung eines durch Doppelkonsonanz gekennzeichneten Kurzvokals
	Syntax	Vokalquantität	Morphologie		Morphologie	Phonem-Graphem- Korrespondenz		Vokalquantität
OLFA	01 Klein- für Großschreibung	09 Einfache Vokalschreibung für markierte Länge	29 Konsonantenzeichen fehlt		19 p t k für b d g im Silbenendrand	33 falscher Konsonant		07 Einfachschreibung für Konsonantenverdopplung
	Gruppe II	Gruppe II	Gruppe I		Gruppe II	Gruppe I		Gruppe II
	alphabetisch	alphabetisch	protoalphabetisch		alphabetisch	protoalphabetisch		alphabetisch

Tabelle 4: Vergleich der Fehlertypologien in AFRA und OLFA anhand der Fehlschreibung *<faratchlos> (*Fahrradschloss*)

Die Gegenüberstellung weist deutlich darauf hin, dass der größte Unterschied zwischen den beiden Fehlerkategorisierungen in der morphologischen Schreibung liegt. Während in der AFRA explizit das fehlende <r> am Morphemanschluss zwischen {fahr} und {rad} als „Konsonantische Ableitung (KA)“ auf der morphologischen Ebene gekennzeichnet wird, ordnet die OLFA diesen Fehler rein deskriptiv als Auslassung eines Konsonantengraphems ein. Allerdings ist eine solche Konsonantenauslassung

an Morphemgrenzen grundsätzlich von der einfachen Tilgung eines Konsonantenzeichens zu unterscheiden – wie etwa bei der Fehlschreibung *<Stuktur> für *Struktur*, bei der keine Morphemgrenze betroffen ist.

Zwar wird die Konsonantenverdopplung in beiden Fehlertypologien berücksichtigt, jedoch differenziert keines der beiden Verfahren zwischen einer silbischen Gelenkschreibung wie <mm> in *kommen* und einer morphologischen, vererbten Schreibung wie <mm> in *kommt*. Dass die Durchdringung des silbischen und des morphologischen Prinzips unterschiedliche Entwicklungsstufen im Orthographieerwerb darstellen, ist in der fachdidaktischen und linguistischen Forschung unbestritten (vgl. u. a. Bredel et al., 2011; Fuhrhop, 2021). Folglich könnte gerade die diagnostische Unterscheidung dieser beiden Phänomene weitaus präzisere Einblicke in den Erwerb von orthographischen Prinzipien der Lernenden liefern.

Diese Notwendigkeit einer feineren Analyse gilt umso mehr, da sich der Orthographieerwerb bei DaF-Lernenden im Vergleich zum Erstspracherwerb weder rein stufenförmig noch linear, sondern dynamisch und komplex gestaltet. Ein linguistisch fundiertes Fehleranalysesystem, das diese Vielfalt an allen orthographischen Fehlerstellen (inkl. möglicher Überlappungen) präzise zu erfassen vermag, kann daher wichtige Informationen sowohl über den individuellen Erwerbsprozess als auch für eine angemessene didaktische Förderung liefern.

2.2 Fremdspracherwerbsypothesen und Orthographieerwerb

Aufbauend auf der Darstellung der orthographischen Prinzipien und ihrer Aneignung im Erstspracherwerb stellt sich die zentrale Frage, inwieweit diese Erkenntnisse auf den Erwerb von Deutsch als Fremdsprache übertragbar sind. Wie bereits deutlich gemacht wurde, unterscheidet sich der Lernprozess von multilingualen Lernenden fundamental von dem der monolingualen Kinder, unter anderem durch bereits vorhandene Alphabetisierungserfahrungen und den Einfluss der erworbenen Sprachen.

Um die spezifischen Bedingungen und Prozesse des orthographischen Erwerbs im DaF-Kontext theoretisch zu fundieren, gibt dieser Abschnitt daher einen Überblick über grundlegende Hypothesen des Fremdsprachenerwerbs. Jede Theorie wird hierbei hinsichtlich ihrer Relevanz für den Erwerb orthographischer Kompetenz kritisch betrachtet.

Zu diesem Zweck werden zunächst die Grundkonzepte der Kontrastivhypothese, der Interlanguage-Hypothese, der Identitätshypothese, des Multi-Kompetenz-Ansatzes sowie der Theorie komplexer dynamischer Systeme (*Complex Dynamic Systems Theory*, CDST) dargestellt. Eine detaillierte Gegenüberstellung dieser Hypothesen wird bezüglich der Fehlerperspektive in der Lernaltersprache, Einflussfaktoren im Erwerbsprozess und der Rolle des Sprachvergleichs ausgelegt. Anschließend werden die Erwerbsypothesen evaluiert, welche als theoretischer Rahmen für die Analyse der vorliegenden Arbeit fungieren.

2.2.1 Von der Kontrastivhypothese bis zur Multi-Kompetenz-Hypothese und Theorie komplexer dynamischer Systeme (CDST)

Zu den frühen Hypothesen der Spracherwerbsforschung zählen die Kontrastivhypothese und die Identitätshypothese. Die kontrastive Linguistik entwickelte sich in den 1940er- und 1950er-Jahren im Kontext behavioristischer Ansätze, insbesondere in der Englischdidaktik, auf Grundlage der Arbeiten von Charles C. Fries (1945) und Robert Lado (1957). Zentrale Annahme war, dass der Kontrast zwischen der Erstsprache (L1) und der Zielsprache (L2) entscheidend für den Lernprozess sei. Lado (1957) betonte die Notwendigkeit eines systematischen Sprachvergleichs zwischen Ausgangs- und Zielsprache, um potenzielle Schwierigkeiten im Fremdspracherwerb vorherzusagen.

Gemäß Lado (1957, S. 2) führen Unterschiede zwischen den beiden Sprachen zu Lernaltschwernissen, während Ähnlichkeiten den Erwerbsprozess erleichtern können. Dieses Prinzip bildet den Kern der Kontrastivhypothese. In den 1950er- und 1960er-Jahren wurde diese weiterentwickelt und ging nun davon aus, dass Lernende sprachliche Strukturen und Muster systematisch aus ihrer L1 auf die L2 übertragen. Strukturelle Übereinstimmungen zwischen L1 und L2 führen dabei zu positivem Transfer, der den Lernprozess erleichtert (Lernerleichterungen), während strukturelle Unterschiede negativen Transfer (Interferenz) verursachen, der zu Lernschwierigkeiten und Fehlern führt. In diesem Sinne gelten die Strukturunterschiede als ausschließliche Ursachen für Schwierigkeiten und Fehler in der Fremdsprache (vgl. Lado, 1961, S. 146, 1967, S. 80). Bausch und Kasper (1979, S. 6) fassten diese starke Version der Kontrastivhypothese in den beiden folgenden Gleichungen zusammen:

- Strukturidentität Grundsprache – Zweitsprache = positiver Transfer = Lernerleichterung = korrekte zweitsprachliche Äußerung
- Strukturdivergenz Grundsprache – Zweitsprache = Interferenz (negativer Transfer) = Lernschwierigkeit = fehlerhafte zweitsprachliche Äußerung

Die Kontrastivhypothese postuliert in ihrer Grundannahme, dass der Orthographieerwerb maßgeblich von den strukturellen Ähnlichkeiten bzw. Unterschieden zwischen der orthographischen Systematik der L1 und der L2 beeinflusst wird. Strukturelle Übereinstimmungen gelten als potenzielle Lernerleichterungen, während Diskrepanzen zu systematischen Fehlern führen können. Demnach lassen sich orthographische Fehler vor allem als Ergebnis negativer Transferprozesse deuten. Das bedeutet: Wenn in der Erstsprache ein bestimmtes orthographisches Prinzip nicht existiert, fehlen den Lernenden entsprechende Anknüpfungspunkte, was die Aneignung in der Fremdsprache erschwert. Vergleicht man beispielsweise zwei DaF-Lerngruppen mit unterschiedlichen L1 wie z. B. Chinesisch und Englisch, würde die Kontrastivhypothese prognostizieren, dass insbesondere für chinesische Lernende grundlegende Schwierigkeiten im Bereich der Graphem-Phonem-Korrespondenz auftreten, da Chinesisch kein alphabetisches Schriftsystem besitzt. Für die Lernenden mit Englisch als L1 sollte hingegen die geteilte alphabetische Schriftbasis sowie die strukturelle Verwandtschaft beider Sprachen – etwa im Bereich ähnlicher Laut-Buchstaben-Zuordnungen oder verwandter Morpheme – eine Erleichterung darstellen. Jedoch zeigen empirische Befunde aus der Fehlerana-

lyse, dass Fehler nicht nur bei strukturellen Differenzen, sondern auch bei vermeintlichen Ähnlichkeiten entstehen können (vgl. Huneke & Steinig, 2013, S. 34). Kritisch anzumerken ist zudem, dass insbesondere die starken Annahmen der Kontrastivhypothese interne Entwicklungsprozesse, individuelle Lernstrategien sowie kontextuelle Einflussfaktoren weitgehend ausblendeten. Diese rein kontrastive Sichtweise wurde später durch differenziertere Ansätze wie die Interlanguage-Hypothese ergänzt, die genau diese Faktoren in den Mittelpunkt stellen.

Die durch nativistische und kognitivistische Annahmen geprägte Identitätshypothese wurde bereits in den 1970er-Jahren als Gegenposition zur kontrastiv geprägten Sichtweise formuliert (vgl. Clahsen, 1982; Meisel, Clahsen & Pienemann, 1981). Sie geht davon aus, dass der Erst- und Fremdspracherwerb im Wesentlichen denselben kognitiven Mechanismen unterliegen. Zentral ist hierbei die Annahme einer angeborenen universellen Grammatik (Universal Grammar, UG), die allen Menschen zur Verfügung steht und den Erwerb beliebiger Sprachen ermöglicht und steuert (vgl. Chomsky, 1965, 1985; Meisel, 2012). Aus dieser Annahme ergibt sich die Folgerung, dass Lernende unabhängig von ihrer Erstsprache eine vergleichbare Entwicklungsabfolge durchlaufen. Empirisch gestützt wurde diese Sichtweise durch frühe Studien von Dulay und Burt (1973, 1974) sowie Bailey, Madden und Krashen (1974), die eine einheitliche Erwerbsreihenfolge im Bereich der englischen Syntax als L2 postulierten (vgl. Hulstijn, Ellis Rod & Søren, 2015, S. 2). Laut der Identitätshypothese sind Fehler der Unvollkommenheit der Kompetenz der Lernenden zuzuschreiben. Die empirische Forschung von Dulay und Burt (1974) beweist, dass nur ein geringer Teil der Fehler in der Lerner Sprache auf den Sprachkontrast zurückzuführen ist. Daher seien die Fehler eher aus der Struktur der Fremdsprache zu erklären. Im Bereich des Orthographieerwerbs würde dies bedeuten, dass Lernende unabhängig von ihrer Erstsprache ähnliche Entwicklungsstufen (siehe Stufenmodelle im Kapitel 2.1.2) durchlaufen, wie im Erstspracherwerb. Die orthographischen Fehler lassen sich somit durch die Komplexität des deutschen Orthographiesystems erklären, wobei ähnliche Fehlererscheinungen bei Lernenden mit unterschiedlichen L1 zu beobachten sind.

Die Identitätshypothese liefert eine sprachvergleichsunabhängige Erklärungsperspektive, vernachlässigt jedoch den Einfluss bereits erworbener Sprachsysteme auf den Erwerb der zweiten oder weiteren Sprache, was den Regeln der Plausibilität und der Lern- und Entwicklungspsychologie widerspricht (vgl. Roche, 2013, S. 120). Andere Einflussfaktoren wie Motivation oder Lernkontext werden in der Kernannahme der Identitätshypothese weniger beleuchtet, da der Fokus stark auf den internen, kognitiven Mechanismen liegt.

Ausgehend von der Kontrastivhypothese entwickelte sich mit der Interlanguage-Hypothese von Selinker (1972) ein paradigmatischer Wandel im Verständnis von Lerner Sprache (vgl. Selinker, 1972, S. 209–210). Im Zentrum dieses Ansatzes steht die Auffassung, dass eine Lerner Sprache ein eigenständiges linguistisches System darstellt, das sich in einem linguistisch-kognitiven Raum zwischen Ausgangs- und Zielsprache entfaltet (vgl. Selinker, 2014, S. 223). Jeder Abschnitt im Erwerbsprozess lässt sich als Zwischenstufe auf dem Weg zur Zielsprache auffassen, die sich durch lernerspezifische

sche Strategien sowie durch systemeigene sprachliche Strukturen in verschiedenen Lernervarietäten auszeichnet (vgl. Glück & Rödel, 2016; Jung & Günther, 2016; Klein, 2000). Gemäß der Interlanguage-Hypothese werden Fehler in der Lernersprache nicht mehr hauptsächlich als Ergebnis von Interferenzen aus der Muttersprache angesehen, sondern vielmehr als wichtiger Bestandteil des Lernprozesses betrachtet (vgl. Corder, 1986, S. 38). Corder (1986, S. 10) unterscheidet zudem zwischen den Begriffen *mistakes* und *errors*: Systematische Fehler (*errors*) geben den Aufschluss über das aktuelle Sprachsystem der Lernenden, während unsystematische Ausrutscher (*mistakes*) auf Performanzprobleme zurückzuführen sind. Dazu führt Selinker das wichtige Konzept der Fossilierung (Selinker, 1972) ein. Mit Fossilierung lässt sich der Umstand beschreiben, dass bestimmte nicht-zielsprachliche Strukturen sich in Lernersprachen verfestigen und im weiteren Erwerbsprozess nur schwer veränderbar sind (vgl. Selinker, 1972, S. 215–216). Dieses „Steckenbleiben“ stellt laut Selinker einen der grundlegenden Unterschiede zwischen Erst- und Zweitspracherwerb dar (vgl. Selinker, 2014, S. 226).

Im Hinblick auf den Orthographieerwerb lässt sich aus Sicht der Interlanguage-Hypothese eine sogenannte „Lerner-Schriftsprache“ erkennen, die orthographische Regelmäßigkeiten der Erstsprache und Zielsprache sowie eigenständige Eigenschaften beinhaltet. Die Lernerorthographie ist dementsprechend nicht als zufällige Ansammlung von Schreibfehlern zu verstehen, sondern als ein sich entwickelndes System, das bestimmte Übergangsformen hervorbringt. Orthographische Schreibungen in der Lerner-sprache zeigen somit auf, dass die Lernenden aktiv versuchen, eigene orthographische Regeln und Strategien aufzubauen. Die orthographischen Fehler in der Lernersprache lassen sich nach Selinker (vgl. 1972, 215 ff.) durch fünf psycholinguistische Prozesse erklären: Transfer aus der Erstsprache oder aus anderen Sprachen, Transfer aus der Lernumgebung wie z. B. durch ungeeignete Lern- oder Übungsmaterialien zum Orthographieerwerb, eingesetzte Lern- und Kommunikationsstrategien und Übergeneralisierungen der orthographischen Regeln und Strukturen.

Zwar wird in der Interlanguage-Hypothese die Eigenständigkeit und Systemhaftigkeit von Lernervarietäten im Vergleich zur Kontrastiv- und Identitätshypothese differenzierter herausgestellt, dennoch wird kritisch angemerkt, dass ihre empirische Fundierung in mehreren Punkten begrenzt erscheint. Insbesondere Henderson (1985, S. 26) kritisiert, dass die Interlanguage-Hypothese keine falsifizierbaren Aussagen generiere und somit keine überprüfbaren Prognosen über den Verlauf des Zweitspracherwerbs zulasse. Die theoretische Annahme eines eigenständigen Zwischensystems zwischen Ausgangs- und Zielsprache bleibe weitgehend spekulativ, da beobachtbare Daten weder eindeutig einer bestimmten Erwerbsphase noch einem spezifischen lernersprachlichen Prozess zugeordnet werden können (vgl. Jordan, 2004; Richards, 1974, 174, 177). Ein weiterer Kritikpunkt an der Interlanguage-Hypothese betrifft ihre begrenzte Berücksichtigung sozialer und kontextueller Faktoren im Zweitspracherwerb. Nach Firth und Wagner (1997) ist die traditionelle Spracherwerbsforschung zu stark auf kognitive Prozesse und individuelle Kompetenzen fokussiert und vernachlässigt dabei die soziale Dimension von Sprachverwendung, wobei weitere Faktoren wie Kommunikationsabsichten, Gesprächsrollen oder situative Anforderungen ebenfalls berücksichtigt werden sollen.

Als eine Ergänzung zum Begriff „Interlanguage“ wurde „Multi-Competence“ erstmals im Kontext des Zweitspracherwerbs verwendet, um „the compound state of a mind with two grammars“ (Cook, 1992) zu bezeichnen (vgl. Cook, 2016, S. 2). Das Konzept der Multi-Kompetenz betrachtet den Zweitspracherwerb als einen Prozess, der den gesamten kognitiven Apparat der L2-Nutzenden einbezieht – nicht nur die Zweitsprache selbst (vgl. Cook, 1992). Die Definition von Multi-Kompetenz wurde später erweitert und führte schließlich zu einer Präzisierung der ursprünglichen Konzeption: „the knowledge of more than one language in the same mind or the same community“ (Cook, 2012). Mit dieser Definition wurde deutlich gemacht, dass sich Multi-Kompetenz nicht auf Syntax beschränkt, sondern auch den Wortschatz, die Phonologie und andere sprachliche Wissensbereiche umfasst; das Konzept wurde über die Zweitsprache hinaus auf weitere Sprachen ausgeweitet; der Geltungsbereich wurde auf ein soziologisches Konstrukt der „Multi-Competence of the community“ erweitert. Dabei wird die sprachliche Vielfalt innerhalb einer Gemeinschaft als kohärentes Ganzes verstanden, nicht als Ansammlung isolierter Sprachsysteme (Cook, 2016, S. 2–3). Entsprechend wird vor allem „L2-user“ statt „L2-learner“ oder „bilingual“ im Multi-Kompetenz-Konzept benutzt (vgl. Cook, 2012). Die L2-Nutzenden haben eine spezifische kognitive Konfiguration; ihre sprachliche Entwicklung steht in einem dynamischen Wechselverhältnis zwischen bereits vorhandenen und neu erworbenen sprachlichen Ressourcen (vgl. Cook, 2016, S. 4). Der Begriff „L2-Lernende“ kann dann benutzt werden, wenn eine Person zwar eine Zweitsprache (L2) lernt, diese aber nicht aktiv im Alltag verwendet (vgl. Cook, 2016, S. 4). Als Beispiel hierfür gelten Germanistikstudent*innen in China, deren primärer Lernkontext der universitäre Unterricht und nicht die alltägliche Kommunikation im zielsprachigen Umfeld ist. Im Gegensatz dazu stehen „L2-Nutzende“, die die Sprache aktiv gebrauchen. Diese Unterscheidung ist für den Orthographieerwerb von wichtiger Bedeutung, da sie direkt die Art und Menge des Inputs betrifft: L2-Nutzende erhalten durch den ständigen Kontakt mit authentischen Texten einen vergleichsweise höheren Input, der zur Verinnerlichung orthographischer Muster beitragen kann. Bei L2-Lernenden im reinen Klassenraumkontext ist dieser Input hingegen quantitativ stark begrenzt und qualitativ didaktisch gesteuert, was ihren Erwerbsprozess maßgeblich beeinflusst (vgl. Edmondson & House, 2011; Lightbown & Spada, 2013).

Im deutschsprachigen Raum wird der Multikompetenz-Ansatz vor allem im Rahmen der Tertiärsprachenforschung behandelt, die sich mit dem Einfluss zuvor gelernter Sprachen auf den Erwerb weiterer Zielsprachen befasst (vgl. Aronin & Hufeisen, 2009). Die Tertiärsprachendidaktik (vgl. Hufeisen, 2003; Hufeisen & Neuner, 2005) schlägt in diesem Zusammenhang insbesondere die gezielte Mobilisierung rezeptiver sprachlicher Kompetenzen vor.

Aus der Perspektive der Multikompetenz-Hypothese lässt sich der Orthographieerwerb in der Fremdsprache als ein Prozess verstehen, der auf einem dynamischen Zusammenspiel verschiedener sprachlicher Wissenssysteme basiert. Dabei wird davon ausgegangen, dass Lernende über ein multiples Sprach- und Schriftwissen verfügen, das nicht isoliert, sondern modular und vernetzt organisiert ist. Die einzelnen Wissensbereiche – etwa L1, L2 und bereits partiell erworbene L3 – stehen in Wechselwirkung

zueinander und bilden eine gemeinsame kognitive Ressource für den Erwerb orthographischer Strukturen. Hans-Bianchi und Katelhön (2010, S. 33) stellt in der Studie über den Orthographieerwerb italienischer Deutschlernender dieses multiple Sprachwissen graphisch dar:

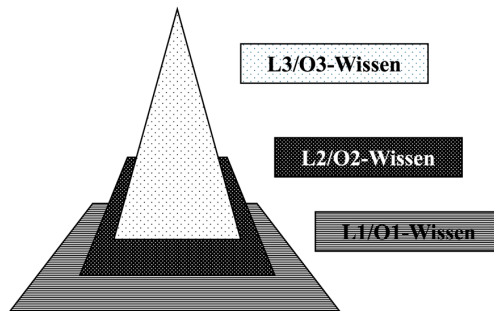


Abbildung 1: Das multiple Sprach- und Orthographiewissen (Quelle: Darstellung in Anlehnung an Hans-Bianchi & Katelhön (2010, S. 33))

In diesem Modell wird das Sprachsystem der L1 als Grundlage verstanden, das den Zugang zu weiteren Sprachen strukturell und strategisch vorbereitet. „Je größer die – auch subjektiv bemessenen – sprachlichen Ähnlichkeiten, desto stärker der Transfer-Effekt. Ebenso ist die L2 in ihrer Spezifität relevant für den Erwerb der L3.“ (Hans-Bianchi & Katelhön, 2010, S. 33).

Die spezifischen Auswirkungen der Multi-Kompetenz sind dabei das Ergebnis eines komplexen Zusammenspiels vielfältiger Faktoren. Hierzu zählen insbesondere individuelle Unterschiede und Erfahrungen im Spracherwerb sowie in der Sprachverwendung. Außerdem sind die Art der Sprachnutzung sowie der spezifische Kontext, in dem die Sprachen erworben und angewendet werden, ebenfalls relevant (vgl. Cook, 2016, S. 4–9).

Während der Multi-Kompetenz-Ansatz die Komplexität des multilingualen Wissenssystems im Individuum sowie innerhalb der Gemeinschaft und die dynamische Interaktion zwischen verschiedenen Sprachen hervorhebt, bietet *Complexity and Dynamic Systems Theory* (CDST) eine breitere theoretische Perspektive, die den gesamten Prozess des Fremdspracherwerbs als ein komplexes, sich selbst organisierendes und dynamisches System betrachtet (vgl. de Bot, 2017; Larsen-Freeman, 1997).

Die CDST stammt ursprünglich aus der Bezeichnung *Dynamic Systems Theory* (DST) der Mathematik und wird seit Mitte der 1990er-Jahre im Diskurs der Angewandten Linguistik verwendet (vgl. Bülow, de Bot & Hilton, 2017, S. 47). Wegweisend waren dabei insbesondere die Arbeiten von Larsen-Freeman (1997), Herdina und Jessner (2002) sowie de Bot, Lowie und Verspoor (2007). Anstelle traditioneller statistischer Methoden in der Zweitspracherwerbsforschung, die primär auf Gruppenmittelwerte und Gesamtverläufe der Lernergruppe fokussieren, rückt die DST die nicht linearen Entwicklungskurven individueller Sprachlernprozesse in den Vordergrund (vgl. Bülow et al., 2017, S. 47).

Der Begriff CDST wurde später von de Bot (2017) eingeführt. Im Bereich der Angewandten Linguistik werden die Begriffe „Complex Systems“ und „Dynamic Systems Theory“ oft synonym verwendet. Beide Konzepte basieren auf der Annahme, dass Sprache und Sprachlernen komplexe, dynamische Systeme sind, die sich durch das Zusammenspiel von Subsystemen und Komponenten verändern und selbst anpassen (vgl. Verspoor & Lowie, 2021, S. 189). Ein zentrales Phänomen in solchen Systemen ist die *Emergenz*. Emergenz beschreibt die Herausbildung qualitativ neuer Eigenschaften, Strukturen oder Verhaltensweisen auf der Makroebene eines Systems, die nicht allein durch die Eigenschaften der einzelnen Komponenten erklärt werden können, sondern aus deren komplexem, nicht linearem Zusammenspiel spontan entstehen (vgl. N. C. Ellis & Larsen-Freeman, 2006; Holland, 2000).

Heute stellt CDST einen leistungsfähigen theoretischen Rahmen dar, um Entwicklungsprozesse im Zweitspracherwerb zu analysieren. Dabei wird betont, dass sprachliche Entwicklung ein individueller, nicht vorhersehbarer und grundsätzlich nicht linearer Prozess ist (vgl. Larsen-Freeman, 1997, S. 152; Verspoor et al., 2021, S. 173).

Während der Begriff *Transfer* im Multi-Kompetenz-Ansatz noch von wichtiger Bedeutung ist (vgl. Cook, 2016, S. 10; Hufeisen, 2003, S. 98–99), wird dieser Begriff von CDST vielmehr als Transformation verstanden: „Transfer is never exact; what is being ‚transferred‘ is reworked to suit a new context“ (Larsen-Freeman, 2017, S. 26). Dies impliziert, dass sprachliches Vorwissen in der CDST nicht als statischer Einflussfaktor verstanden wird, sondern als dynamisch umstrukturierte Ressource, die sich kontextabhängig neu organisiert.

Ein Schlüsselkonzept in Untersuchungen zu CDST ist Variabilität (*variability*). Insbesondere die zwischen Lernenden beobachteten Unterschiede werden als interindividuelle Variabilität betrachtet. Die bei einzelnen Lernenden beobachteten Entwicklungsunterschiede sind als intraindividuelle Variabilität zu interpretieren. Variabilität wird zudem als eine intrinsische Eigenschaft eines sich selbst organisierenden und entwickelnden Systems verstanden, in dem die Lernenden ihre eigenen Entwicklungspfade verfolgen (vgl. de Bot & Makoni, 2005). Solche Entwicklungspfade können durch sogenannte Attraktorzustände (*Attractor States*) gekennzeichnet sein. Attraktorzustände sind relativ stabile Phasen oder Muster im Systemverhalten, zu denen das System immer wieder tendiert, und stellen eine temporäre Stabilität im Entwicklungsprozess dar (Larsen-Freeman & Cameron, 2008, 59 ff.; vgl. z. B. van Geert, 1991). Die Analyse von inter- und intraindividuelle Variabilität sowie Attraktorzuständen könnte daher Hinweise darauf geben, wann und warum Veränderungen in der zweitsprachlichen Entwicklung von Lernenden stattfinden (vgl. Schwendemann, 2023, S. 81).

Komplexe dynamische Systeme verändern sich unter dem Einfluss vielfältiger interner und externer Faktoren über die Zeit. Dieser Zusammenhang wird von de Bot und Larsen-Freeman (2011, S. 11) grafisch wie folgt dargestellt:

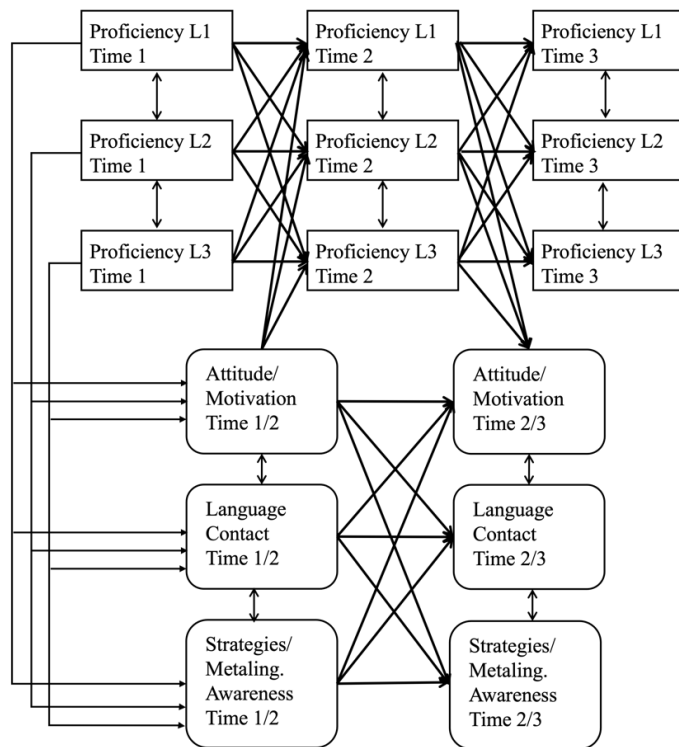


Abbildung 2: Zusammenspiel aller Faktoren über Zeit (Quelle: de Bot/Larsen-Freeman 2011: 11)

Nach dieser Darstellung ist die sprachliche Kompetenz eines Lernenden zu jedem Zeitpunkt von einer Vielzahl miteinander verknüpfter Faktoren abhängig. Gleichzeitig spielen die weiteren Faktoren wie z. B. Lernmotivation, Einstellung zum Gegenstand, persönliche Umstände, das vorhandene Vorwissen über andere Sprachen, bekannte Lernstrategien, Alter, sozial-kontextuelle Faktoren usw. zu allen Zeitpunkten eine Rolle.

Übertragen auf den Orthographieerwerb bedeutet die CDST, dass die Entwicklung der Rechtschreibkompetenz als ein dynamisches, komplexes System verstanden wird. In diesem Kontext können Attraktorzustände, in Anlehnung an Bulut (2018, S. 53), als die üblicherweise in Rechtschreibentwicklungsmodellen definierten Stufen interpretiert werden. Nach Bulut (2018, S. 53) kann diese Interpretation erklären, weshalb die Rechtschreibentwicklung zum einen durch Stabilität gekennzeichnet sein kann, zum anderen aber auch dynamisch und individuell fortschreitet. Dieses komplexe Orthographiesystem besteht nicht allein aus dem Wissen über orthographische Regularitäten, sondern vielmehr aus der komplexen Interaktion dieses Wissens mit anderen sprachlichen Subsystemen (wie Phonologie, Morphologie, Syntax in L1, L2, L3, ...), kognitiven Prozessen, angewandten Lernstrategien sowie dem sprachlichen Input (vgl. Verspoor & Lowie, 2021, S. 189). Darüber hinaus wird es durch eine Vielzahl interner und externer Einflussfaktoren moduliert (vgl. Larsen-Freeman, 1997, S. 151), zu denen beispielsweise Motivation, der spezifische Lernkontext, die Qualität des Unterrichts, die Menge der

Schreibpraxis und auch der L1-Hintergrund des Lernenden zählen (vgl. Verspoor & Lowie, 2021, S. 190). Die beobachtbare Rechtschreibleistung zu einem gegebenen Zeitpunkt ist somit als das emergente Ergebnis der komplexen, prinzipiell nicht linearen und unvorhersehbaren Wechselwirkungen all dieser Elemente im individuellen System des Lernenden zu verstehen (vgl. Larsen-Freeman, 1997, S. 152; Verspoor et al., 2021, S. 173). Genau an dieser umfassenden und hochgradig abstrakten Modellierung setzt ein zentraler Kritikpunkt an (vgl. Gregg, 2010): Durch die Grundannahme, dass potenziell alle internen und externen Faktoren nicht linear interagieren, entzieht sich die CDST in weiten Teilen einer strengen empirischen Überprüfbarkeit. Spezifische Kausalitäten lassen sich methodisch kaum noch isolieren, wodurch die Formulierung überprüfbarer Gegenargumente massiv erschwert wird (vgl. ebd., S. 554, 557). Trotz dieser empirischen Herausforderungen bietet das Modell jedoch einen bedeutenden Mehrwert. Für den Orthographieerwerb kann diese Emergenz, beispielsweise in Form der Ausbildung einer neuen Regel oder eines neuen Musters im Schreibverhalten, auch als innerer Regelbildungsprozess verstanden werden (vgl. Bulut, 2018, S. 51).

Zusammenfassend zeigt die Entwicklungslinie der Fremdspracherwerbsforschung eine zunehmende Ausdifferenzierung der Erklärungsmodelle. Während frühe Ansätze wie die Kontrastiv- und Identitätshypothese grundlegende Mechanismen wie Transfer und universelle Erwerbsstufen isolierten, offenbarte bereits die Interlanguage-Hypothese die Notwendigkeit, Lernersprache als eigenständiges, dynamisches System zu begreifen. Dennoch stoßen auch diese Ansätze an ihre Grenzen: Sie unterschätzen entweder die systemische Verflechtung der L1 mit weiteren Fremdsprachen, oder sie können die hochgradig variable und nicht lineare Natur individueller Entwicklungsverläufe nicht adäquat modellieren. Es entsteht somit eine Erklärungslücke, die ein Framework erfordert, das sowohl die kognitive Realität der Mehrsprachigkeit als auch die komplexen Prozesse der Entwicklung fassen kann. Das folgende Kapitel wird daher begründen, warum ein integrierter Bezugsrahmen, der die Multi-Kompetenz-Hypothese mit der Complex Dynamic Systems Theory (CDST) synthetisiert, dieser doppelten Anforderung gerecht wird und somit die theoretische Rahmung der vorliegenden Untersuchung darstellt.

2.2.2 Begründung der theoretischen Rahmung

In der vorliegenden Untersuchung wird die orthographische Kompetenz chinesischer Deutschlernender untersucht. Die untersuchten Lerngruppen haben die deutsche Sprache im Fremdsprachenunterricht gelernt, wobei Mehrsprachigkeit ein gegebener Bestandteil ist. Zhao (2020, S. 51) beschreibt eine aktuelle Situation von Deutsch als Fremdsprache in China und fasst das Lernen der deutschen Sprache in China in vier Formen zusammen:

- a) im vierjährigen Hauptfach- bzw. Germanistikstudium,
- b) als studienbegleitende DaF-Kurse,
- c) als Schulfach in der Mittelschule,
- d) als kommerzielle Kursangebote an öffentlichen oder privaten Schulen.

Die Gruppen a) und b) bilden die Hauptgruppe der Deutschlernenden in China. Obwohl seit 2018 Deutsch neben Japanisch, Französisch und Spanisch als erste Fremdsprache in der Schule eingeführt und dementsprechend auch als Fach der staatlichen Hochschuleaufnahmepprüfung angeboten wird, wird Deutsch als Schulfach im Bereich der Mittelschule eher in den wirtschaftlich entwickelten Städten wie Shanghai, Beijing, Shenzhen unterrichtet. Außerdem spielt die Gruppe d) nur eine ergänzende Funktion für Deutsch als Fremdsprache in China (vgl. Zhao, 2020). Die vorliegende Arbeit konzentriert sich auf die Hauptgruppe a), wobei die meisten Germanistikstudierenden bis zum Studienanfang bereits Englisch als erste Fremdsprache für sechs Jahre in der Schulbildung lernten.

Angesichts des mehrsprachigen Hintergrunds der chinesischen Deutschlernenden bildet die Multi-Kompetenz-Hypothese den ersten unverzichtbaren Pfeiler des Bezugsrahmens. Als mehrsprachiges Individuum verfügen die chinesischen Lernenden über sprachliches Wissen, Lernerfahrungen sowie bereits entwickelte Lernstrategien (vgl. Hufeisen & Neuner, 2005) sowohl in ihrer Erstsprache Chinesisch als auch in der ersten Fremdsprache Englisch. Sie allein ermöglicht es, die Lernenden nicht als defizitäre Anfänger, sondern als kompetente „L2-user“ mit einem komplexen, vernetzten Sprachwissen zu konzeptualisieren. Um den Einfluss des sprachlichen Vorwissens aus der Erst- und den weiteren Fremdsprachen auf die deutsche Orthographie tiefgreifend zu analysieren, liefert dieses Modell die theoretische Legitimation. Phänomene werden so nicht als simple „Interferenzen“ abgetan, sondern als logische Ergebnisse der Interaktion innerhalb eines einzigen, multikompetenten kognitiven Systems.

Konkret auf den Orthographieerwerb bezogen bedeutet dies, dass Lernende über ein multiples Sprach- und Schriftwissen verfügen, das nicht isoliert, sondern modular und vernetzt organisiert ist (vgl. Hans-Bianchi & Katelhön, 2010). Die orthographischen Kenntnisse aus dem Chinesischen (z. B. visuell-motorische Strategien, *Pinyin*-Graphem-Phonem-Zuordnung), dem Englischen (z. B. Graphem-Phonem-Korrespondenzen) und dem Deutschen stehen in ständiger Wechselwirkung und bilden eine gemeinsame kognitive Ressource. Dieser Ansatz erlaubt es, Schreibphänomene als Ergebnis der Aktivierung von Wissen aus verschiedenen sprachlichen Teilrepertoires zu deuten.

Während die Multi-Kompetenz-Hypothese primär den Zustand und die Struktur des Wissenssystems beschreibt, liefert sie allein kein ausreichendes Instrumentarium zur Analyse seiner Entwicklung über die Zeit. Um den Verlauf der Kompetenzentwicklung, das Auftreten von Variabilität und die nicht linearen Lernpfade zu modellieren, ist ein prozessorientiertes Modell notwendig. Genau diese Lücke schließt die Complex Dynamic Systems Theory (CDST) als zweiter Pfeiler des Rahmens. Sie ermöglicht es, die Lernautsprache als komplexes, sich selbst organisierendes System zu verstehen, dessen Entwicklung durch inter- und intraindividuelle Variabilität gekennzeichnet ist.

Die Anwendung der CDST auf den Orthographieerwerb bedeutet, diesen als dynamischen Prozess zu verstehen, der nicht rigide, linear ablaufende Stufenmodelle (vgl. u. a. May, 1990; Scheerer-Neumann, 1987; Thomé, 2006; Valtin, 1994) durchläuft. Vielmehr können die in diesen Modellen definierten Stufen als temporär stabile Attraktorzustände interpretiert werden, in denen das System der Lernenden für eine ge-

wisse Zeit verweilt (vgl. Bulut, 2018, 52 f.). Die Rechtschreibleistung zu einem gegebenen Zeitpunkt ist somit das emergente Ergebnis der komplexen Interaktion aller sprachlichen Subsysteme (Phonologie, Morphologie, Syntax etc. aus L1, L2 und Lx), kognitiver Ressourcen und externer Faktoren. Fortschritt entsteht hier durch die Destabilisierung eines bestehenden Zustands und die anschließende Reorganisation des Systems auf einer neuen, komplexeren Ebene.

Die theoretische Erklärungskraft des hier gewählten Bezugsrahmens gründet auf der systematischen Integration zweier komplementärer Ansätze: der CDST und der Multi-Kompetenz-Hypothese. Während die Multi-Kompetenz-Hypothese die gleichsam statische Strukturebene des Spracherwerbs konzeptualisiert, indem sie die kognitive Architektur des mehrsprachigen Individuums und dessen vernetzte Wissenssysteme postuliert, fokussiert die CDST auf die Prozessebene. Sie modelliert die dynamischen Mechanismen, die der Entwicklung und Reorganisation dieses kognitiven Systems im Zeitverlauf zugrunde liegen. Erst die Verschränkung dieser strukturellen Stufenabfolgen und prozessorientierten Fassung schafft die Voraussetzung für eine vielfältige Analyse des Orthographieverwerbs. Orthographische Phänomene lassen sich demzufolge nicht nur auf ihre Quellen im multikompetenten Wissenssystem zurückführen, sondern auch als emergente Ergebnisse von Entwicklungsdynamiken, Variabilität und temporär stabilen Phasen interpretieren.

Ein zentraler Begriff für die Analyse von sprachenübergreifenden Einflüssen ist der des Transfers. Während die CDST diesen Prozess treffend als *Transformation* (Larsen-Freeman, 2017, S. 26) beschreibt, um die dynamische und niemals identische Übernahme von Strukturen zu betonen, wird in der vorliegenden Arbeit für die analytische Praxis bewusst der etablierte Begriff *Transfer* beibehalten. Insbesondere im Rahmen der Multi-Kompetenz-Hypothese ist es für die lernerkorpusbasierte Fehleranalyse besonders sinnvoll, die mögliche Einflussrichtung von einer identifizierbaren Quelle (der L1 Chinesisch oder der L2 Englisch) auf die Zielsprache Deutsch zu benennen. Der Transferbegriff wird hierbei jedoch im Lichte beider Theorien neu gefasst: Er bezeichnet den Prozess, bei dem orthographisches Wissen aus einem bereits existierenden Sprachsystem aktiviert und für die Zielsprache nutzbar gemacht wird, wobei dieser Prozess stets eine Umstrukturierung und Anpassung an den neuen sprachlichen Kontext impliziert.

Für den zentralen Forschungsgegenstand der vorliegenden Arbeit wird der Begriff „Fehler“ beibehalten, um abweichende Rechtschreibphänomene in der Lersprache zu analysieren, jedoch nicht im Sinne eines defizitären Verständnisses verwendet. Die beiden Begriffe „Fehler“ und „Abweichung“ werden funktional identisch gebraucht. Fehler werden im Sinne dieser Untersuchung als entwicklungsbezogene Phänomene verstanden, die durch individuelle Lernwege, systeminterne Umstrukturierungen und temporäre Instabilitäten bedingt sind. Obwohl es durchaus schwierig ist, Ursache-Wirkungs-Beziehungen zwischen Variabilität und Veränderung zu erfassen und zu interpretieren und diese häufig auch multikausal sind (vgl. van Dijk, Verspoor & Lowie, 2011, S. 56–57), können Unterschiede im Grad der Variabilität Hinweise auf Entwicklungsprozesse verdeutlichen (vgl. Spoelman & Verspoor, 2010, S. 532). Dabei wird angenom-

men, dass Variabilität die Ursache für eine Entwicklung ist und als Indikator für einen bestimmten Zeitpunkt innerhalb dieser Entwicklung angesehen werden kann (vgl. de Bot et al., 2007). Die Analyse der orthographischen Kompetenz erfolgt somit unter der Annahme, dass orthographische Fehler als eine Facette der Variabilität wichtige Einblicke in die Dynamik des Lernens bieten. Orthographische Fehler markieren in diesem Verständnis mögliche Attraktorzustände (vgl. de Bot & Larsen-Freeman, 2011, S. 14) oder Übergangszonen innerhalb individueller Lernpfade.

Zusammenfassend erlaubt die Kombination der beiden Ansätze, orthographische Kompetenz chinesischer Deutschlernender als Ergebnis eines dynamischen Entwicklungsprozesses zu betrachten. Bei der Fehleranalyse wird die Multikausalität sowie das komplexe Zusammenspiel des multiplen Sprachwissens mit anderen internen und externen Einflussfaktoren berücksichtigt.

2.3 Das sprachliche Repertoire chinesischer Deutschlernender

Die kombinierte theoretische Rahmung aus Multi-Kompetenz-Ansatz und CDST verlangt, die Lernenden nicht als „Deutsch-Anfänger*innen“, sondern als mehrsprachige Individuen zu beschreiben, deren bereits vorhandene Sprachsysteme wechselseitig aufeinander einwirken. Im Folgenden wird daher das sprachliche Repertoire chinesischer Deutschlernender entfaltet, wobei der Fokus auf der Ausgangssprache Chinesisch sowie der Zielsprache Deutsch liegt. Aufbauend auf der Erörterung des multiplen Sprachwissens dieser Lernendengruppe (Kapitel 2.3.1) werden anschließend transferrelevante Aspekte als Einflussfaktoren für die dynamische Interaktion der beteiligten Sprachen analysiert (Kapitel 2.3.2). Den Abschluss bildet, gestützt auf bestehende Studien, eine detaillierte Untersuchung des Transferpotenzials zwischen dem Chinesischen und dem Deutschen speziell auf orthographischer Ebene (Kapitel 2.3.3).

2.3.1 Das multiple Sprachwissen chinesischer Deutschlernender

Die chinesische Sprache wird heute von ca. 1,2 Milliarden Menschen gesprochen und gilt damit als die weltweit am häufigsten gesprochene Erstsprache (vgl. Eberhard, Simons & Robinson, 2026). Sie ist vor allem in der Volksrepublik China beheimatet, wird darüber hinaus aber auch in weiteren ost- und südostasiatischen Ländern gesprochen: Singapur, Indonesien, Malaysia, Taiwan und Thailand. Jedoch trifft der Begriff „Chinesisch“ auf einen Gesamterbegriff, der das Hochchinesische (pǔtōnghuà, Mandarin) und sechs weitere im Süden vertretene chinesische Dialektbündel (Xiang, Kantonese, Wu, Min, Hakka und Gan) umfasst (vgl. Hunold, 2021, S. 1).

Diese Dialekte unterscheiden sich nicht nur in phonologischen, sondern auch in lexikalischen und syntaktischen Merkmalen sehr stark voneinander (vgl. Hunold, 2021, S. 1). Da diese Unterschiede so erheblich sind, ist eine gegenseitige Verständigung zwischen Sprecherinnen und Sprechern verschiedener Regionen – teils sogar innerhalb derselben Provinz – oft nur eingeschränkt möglich. Um diese Kommunikationsbarriere

ren zu überwinden und die nationale Alphabetisierung voranzutreiben, wurde im Oktober 1955 eine Standardisierung der chinesischen Sprache eingeleitet (vgl. T. Liu, 2015, S. 22). Diese Reform war weit mehr als eine rein sprachpolitische Maßnahme; sie markierte den Übergang zu einer modernen, normierten Standardsprache. Die chinesische Sprache wurde vom Bildungsministerium (Staatsrat der VR China 1956) als „Pǔtōnghuà“ bezeichnet und wie folgt definiert:

普通话是以北京语音为标准音，以北方话为基础方言，以典范的现代白话文著作作为语法规范。 *Putonghua shi yi Beijing yuyin wei biao zhun yin, yi beifanghua wei jichu fangyan, yi dianfan de xiandai baihuawen zhuzuo wei yufa guifan.* (Putonghua nimmt die Pekinger Lautung als Standardphonetik, die nördlichen Dialekte als Basisdialekt und die beispielhaften modernen Baihua-Werke als grammatische Norm.)

Seitdem dient das moderne chinesische *Putonghua* als Standardsprache im ganzen Festland-China und zurzeit auch in Hongkong. In der vorliegenden Arbeit bezeichnet *die chinesische Sprache* die Standardsprache *Putonghua*, wohl wissend, dass im Einzelnen individuelle und dialektale Einflüsse vorhanden sind.

Ein systematischer Vergleich der chinesischen und deutschen Orthografie muss von dem fundamentalen Unterschied ihrer Schriftsysteme ausgehen. Das chinesische Schriftsystem ist, im Gegensatz zu den meisten europäischen Sprachen, nicht alphabetisch, sondern logographisch. Es wird durch Schriftzeichen (Sinogramme) realisiert, aus denen sich keine direkten Aussagen zu Phonem-Graphem-Beziehungen ableiten lassen. Das deutsche Schriftsystem basiert hingegen auf dem alphabetischen Prinzip. Aufgrund dieses grundlegenden typologischen Unterschieds sind die beiden Orthographiesysteme nur bedingt direkt vergleichbar. Um einen systematischen Vergleich dennoch zu ermöglichen, wird in der vorliegenden Arbeit die offizielle Umschrift *Pinyin* als gemeinsame Analyse- und Vergleichsgrundlage für das Chinesische herangezogen.

Am 6. Februar 1956 wurde 汉语拼音方案 *Hanyu Pinyin Fang'an* (wörtl. Entwurf einer chinesischen Phonographie) von dem Komitee für die Schriftreform der chinesischen Sprache vorgelegt. Am 11. Februar 1958 wurde dieser Entwurf in der 5. Tagung des 1. Nationalen Volkskongresses genehmigt (vgl. Cremerius, 2012, S. 16; Hinch, 2003, S. 145; M. Wang, 1993, S. 21). Am 31. Oktober 2000 wurde das 中华人民共和国国家通用语言文字法 *Zhonghua Renmin Gongheguo Guojia Tongyong Yuyan Wenzifa* (wörtl. Gesetz über die national allgemein gebräuchliche Sprache und Schrift der VR China) verabschiedet. Der Paragraph § 18 betrifft das *Hanyu Pinyin Fang'an*:

- 国家通用语言文字以《汉语拼音方案》作为拼写和注音工具。 *Guojia Tongyong yuyan wenzi yi "Hanyu Pinyin Fang'an" zuwei pinxie he zhuyin gongju.* (Die national allgemein verwendete Sprache und Schrift benutzen *Hanyu Pinyin Fang'an* als Instrument für die Schreibung bzw. die Angabe der Aussprache.).
- 《汉语拼音方案》是中国人名、地名和中文文献罗马字母拼写法的统一规范，并用于汉字不便或不能使用的领域。 *"Hanyu Pinyin Fang'an" shi zhongguo renming, diming he zhongwen wenxian luoma zimu pinxiefa de tongyi guifan, bing yongyu hanzi bubian huo buneng shiyong de lingyu* („*Hanyu Pinyin Fang'an*“ ist die einheitliche Norm für die Umschrift von chinesischen Namen, Ortsnamen und

die Romanisierung chinesischer Literatur und wird in Bereichen verwendet, in denen chinesische Schriftzeichen unpraktisch oder nicht verwendbar sind.)

- 初等教育应当进行汉语拼音教学 *chudeng jiaoyu yingdang jinxin hanyu pinyin jiaoxue* (*Hanyu pinyin* sollte in der Grundbildung vermittelt werden)

Eines der Hauptziele des 1958 offiziell genehmigten *Hanyu Pinyin Fang'an* war die Förderung der landesweiten Alphabetisierung und die Standardisierung der chinesischen Hochsprache. Das *Pinyin*-System wurde daher zunächst grundlegend im Schulunterricht eingesetzt, um Kindern das Erlernen der Sinogramme zu erleichtern. Darüber hinaus diente es als Umschrift in offiziellen Kontexten wie auf Straßenschildern, in Lexika oder auf Plakaten, oft parallel zu den Schriftzeichen (vgl. Cremerius, 2012, S. 16; M. Wang, 1993, S. 22). Das parallele Erscheinen von Sinogrammen und deren *Pinyin* schafft oft den Eindruck, dass *Pinyin* eine reine wörtliche Transkription für die chinesischen Schriftzeichen sei. Jedoch argumentiert Hinch (2003) in ihrer Studie, dass *Pinyin* in engerem Sinne keine Transkription ist, weil nach *Hanyu Pinyin Fang'an* orthographische Regelungen bei der Schreibung von *Pinyin* verwendet werden (vgl. Hinch, 2003, 150 f.):

- <y> und <w> dienen als künstliche und willkürliche Anlaute vor einigen Vokalen (/i/, /u/, /o/) in Wörtern wie z. B. *yi* (/i/), *yin* (/in/), *ying* (/iŋ/) und *wu* (/u/). Es sind phonetisch keine Konsonanten.
- Der Hauptvokal /o/ im Wort *jiou* /tɕiao/ wird weggelassen und als *jiu* geschrieben.
- Im Alphabet gibt es zu jedem Buchstaben einen Großbuchstaben, der keine zusätzliche phonetische Funktion hat.
- Der Laut [y] wird nach <n> und <l> als *ü*¹ geschrieben, wie z. B. in den Wörtern *nü* (/ny/) oder *lǚ* (/ly/); nach <j>, <q>, <x> und <y> als *u*, wie in den Wörtern *ju* (/tɕy/) und *qu* (/tɕhy/).

Diese Feststellung, dass *Pinyin* als eigene Schrift bezeichnet werden kann, ist laut Hunold (2009) von großer Bedeutung. Allerdings betont sie einschränkend, dass das *Pinyin*-System kein eindeutiges Abbildungssystem von Lautung und Bedeutung darstelle. Ein allein mit *Pinyin* geschriebener Satz sei aufgrund der zahlreichen Homonyme meist schwer verständlich, was dazu führe, dass die semantische Eindeutigkeit ohne die Verwendung von Schriftzeichen verloren gehe (vgl. Hunold, 2009, S. 48).

Studien belegen, dass das Erlernen von *Pinyin* einen positiven Effekt auf den späteren Erwerb der englischen Sprache, insbesondere auf die phonologische Bewusstheit, haben kann (vgl. McBride-Chang et al., 2004; Yin et al., 2011). So zeigen McBride-Chang et al. (2004), dass Kinder in Peking auf der Silben- und phonologischen Ebene deutlich besser abschneiden als ihre Altersgenossen in Hongkong. Dieser Vorteil sei auf ihre Kenntnisse des *Pinyin* zurückzuführen. In ähnlicher Weise verglichen Cheung, Chen,

1 Das *Pinyin*-Graphem <ü> für das Phonem /y/ wurde in der Reform 2012 abgeschafft und stets als <v> geschrieben wie z. B. *lv* (/ly/), *ju* (/tɕy/). Die Lernenden in der vorliegenden Arbeit haben jedoch vor 2012 ihre Grundbildung des *Pinyin*-Systems schon abgeschlossen, deshalb wird in dieser Arbeit für den Laut [y] das *Pinyin*-Graphem <ü> verwendet, wenn Beispiele in *Pinyin* aufgeführt sind.

Lai, Wong und Hills (2001) die phonologische Bewusstheit im Englischen von Kindern mit Kantonesisch als L1 in Guangzhou und Hongkong. Sie wiesen einen signifikanten Vorteil durch das *Pinyin*-Training bei den Kindern in Guangzhou nach, was die Analysen von Anlaut, Reim und Koda betrifft, im Vergleich zu ihren Altersgenossen in Hongkong. Ähnliche Ergebnisse wurden in der Studie von Zhou und McBride (2022, S. 8) über die Transfer-Effekte von phonologischen Fähigkeiten in Englisch, Mandarin und Kantonesisch bei Kindern mit Kantonesisch als L1 gewonnen. Sie veranschaulichen, dass starke phonologische Fähigkeiten in Kantonesisch positiv mit denen in Mandarin und Englisch korrelieren.

Bislang existieren jedoch kaum Studien, die den Einfluss von *Pinyin*-Training auf die phonologische Bewusstheit beim Erwerb von Deutsch als Fremdsprache untersuchen. Viele Forscher*innen sind der Meinung, dass Lernende mit Chinesisch als L1 (Aussprache-)Schwierigkeiten im Bereich der englisch-deutschen Phonem-Graphem-Interferenzen aufweisen, weil Deutsch im chinesischen Bildungskontext meist erst nach dem Englischen als zweite Fremdsprache erworben wird (vgl. Hunold, 2009, S. 49, 2021, S. 8).

Dennoch entstand in den meisten Untersuchungen zu den lernersprachlichen Fehlern eine Übereinstimmung insofern, dass Fehler im Bereich Orthographie auch mit Schwierigkeiten bei der Aussprache im Zusammenhang stehen (siehe Kapitel 2.4). Es wird in vielen Studien als Fazit gezogen, dass die sprachsystemischen Unterschiede zwischen Chinesisch und Deutsch einen großen Einfluss auf die Ausspracheprobleme chinesischer Lernender haben (Forschungsstand zum Aussprachefehler chinesischer Deutschlernender siehe Abschnitt 2.3.3).

Die strukturellen Unterschiede zwischen dem Chinesischen (inkl. *Pinyin*) und dem Deutschen sind aus der Perspektive der Multi-Kompetenz und der Theorie komplexer dynamischer Systeme (CDST) von besonderer Relevanz, da sie potenzielle Spannungsfelder und Interaktionspunkte im kognitiven System der Lernenden darstellen. Obwohl diese Relevanz in der Theorie zwar anerkannt ist, mangelt es bislang an detaillierten kontrastiven Analysen, die sich spezifisch mit dem Orthographieerwerb befassen. Im Folgenden konzentriert sich die Darstellung daher vor allem auf die Vokal- und Konsonantensysteme sowie die Silbenstrukturen des Chinesischen und Deutschen.

Obwohl das englische Sprach- und Schriftsystem unbestritten ebenfalls Teil des sprachlichen Repertoires chinesischer Deutschlernender ist (vgl. dazu auch die Ausführungen in Kapitel 2.2.2), wird in diesem Abschnitt auf einen detaillierten kontrastiven Vergleich unter expliziter Einbeziehung des Englischen (sowohl zwischen Deutsch und Englisch als auch zwischen Chinesisch und Englisch) verzichtet. Eine fundierte Analyse, die das Englische als intervenierende Variable konsequent berücksichtigt, erscheint vor dem Hintergrund der vorliegenden Datengrundlage nur eingeschränkt sinnvoll, da keine differenzierten Informationen zur tatsächlichen Englischkompetenz der Probanden im Prüfungskorpus vorliegen. Dennoch wird an relevanten Stellen punktuell auf Transferphänomene aus dem Englischen hingewiesen, sofern diese dabei helfen, die spezifischen orthographischen Abweichungen im Lernertext besser zu erklären.

2.3.2 Kontrastive Aspekte als Einflussfaktor dynamischer Interaktion

Das Vokalsystem der deutschen Sprache umfasst 16 Vokalphoneme (Monophthonge), die anhand der Kriterien Zungenposition (vertikal: hoch, mittel, tief; horizontal: vorne, zentral, hinten), Gerundetheit und Gespanntheit unterschieden (vgl. Eisenberg, 2013, S. 88; Fuhrhop, 2021) werden können und wie folgt in der Tabelle 5 dargestellt werden.

Tabelle 5: Das Vokalsystem des Deutschen (Quelle: Eisenberg, 2016, S. 36; Hunold, 2009, S. 87)

Zungenlage Zungenhöhe	vorne		zentral	hinten
	ungerundet	gerundet		
hoch	i: ɪ	y: ʏ		u: ʊ
mittel	e: ɛ: ɛ	ø: œ	ə ɐ	o ɔ
tief			a a:	

Ein weiterer Vokalparameter ist die Gespanntheit (gespannt/ungespannt), damit wird der Grad der Muskelspannung der Zunge bei der Artikulation von Vokalen gekennzeichnet (vgl. Meinhold & Sock, 1982, S. 28). Im Deutschen sind lange Vokale in der Regel gespannt und kurze ungespannt. Wenn man Paare von gespannten und ungespannten Vokalen im Deutschen ansetzt (wie etwa /i:/ – /ɪ/), bleiben zwei in der Mitte stehende Vokale /ə/ und /ɐ/ übrig. Sie sind die beiden Reduktionsvokale: das *e*-Schwa und das *a*-Schwa, die in der Regel nur in unbetonten Reduktionssilben zweisilbiger Wörter (wie der jeweils zweiten Silbe in *haben* oder *lieber*) auftreten (vgl. Eisenberg, 2013, S. 67).

Für den Kernwortschatz des Deutschen werden normalerweise drei Diphthonge angenommen: /aʊ/ in *Baum*; /aɪ/ in *Mais*, *Leim*, *Bayern* oder *Speyer* und /ɔɪ/ in *heute* oder *Läufer*. Diese drei Diphthonge gehören zur Klasse schließender Diphthonge, weil der zweite Bestandteil geschlossener als der erste ist. Außerdem gibt es im Deutschen 15 weitere Diphthonge (wie z. B. /i:ɐ/ in *wir*), die auf die vokalisierte Aussprache von *r* zurückgehen. Da das *a*-Schwa als zentraler Vokal gilt, handelt es sich um zentralisierende Diphthonge (vgl. Fuhrhop & Peters, 2013, 27 f.).

Wenn man das Vokalsystem des Chinesischen beschreiben möchte, ist es zunächst wichtig zu erwähnen, dass die chinesische Sprachwissenschaft auf der Ebene der Silbenbeschreibung einer anderen traditionellen Segmentierung folgt als die deutsche. Statt der Begriffe „Vokale“ und „Konsonanten“ wird eine chinesische Silbe strukturell in „Anlaut“ (声母 *shengmu*) und „Auslaut“ (韵母 *yunmu*) unterteilt. Diese strukturellen Positionen werden phonetisch auf der Ebene der Phoneme durch Konsonanten und Vokale realisiert. Der Auslaut im Chinesischen besteht aus einem Vokal (Monophthong), einem Di- oder Triphthong. Der Anlaut besteht immer, abgesehen von Affrikaten, aus einem einzelnen Konsonanten (vgl. T. Liu, 2015, S. 37).

Obwohl in *Hanyu Pinyin Fang'an* phonematisch sechs vokalische Hauptklassen unterschieden und durch die sechs *Pinyin*-Grapheme <a> (/a/), <o> (/o/), <e> (/ɛ/), <i> (/i/), <u> (/u/), <ü> (/y/) repräsentiert werden, ist die Anzahl der Monophthonge im Chinesischen umstritten, in der Literatur variieren die Zahlen zwischen neun und siebzehn. Hunold (2021, S. 4) identifiziert neun Vokalphoneme im Chinesischen und kategorisiert die Vokale nach Zungenhebung (siehe Tabelle 6):

Tabelle 6: Die Vokale des Chinesischen (Quelle: Hunold 2021: 4)

Zungenhebung	Vorderzunge	Mittelzunge	Hinterzunge
Hoch	i y		u
Mittel	ɛ	ə ɐ	ɤ o
Niedrig	a		

Ein finales *-i* (/ɿ/, /ɨ/) nach *z-* (/ts/), *c-* (/tʃ/), *s-* (/s/), *zh-* (/tʃ/), *ch-* (/tʃʰ/), *sh-* (/ʃ/) und *r-* (/ʒ/) wird laut Hunold (2009, S. 87) nicht als eigenständiger Vokal betrachtet, sondern als allophonische Variante von /i/, da dieses finale *-i* sehr schwach ausgesprochen wird.

Chiao und Kelz (1985, S. 125) behaupten hingegen, dass das Hochchinesische 17 Vokalphoneme enthält. Im Vergleich dazu stimmen viele chinesische Linguist:innen (vgl. Huang & Liao, 2017, 101 ff.; Lin & Wang, 2013, 45 ff.; Qian, 2006, S. 13) teilweise mit Hunold überein und meinen, dass die chinesische Sprache nur über wenige Vokale verfügt, sie zählen jedoch das finale *-i* (/ɿ/, /ɨ/) zu den Vokalen, auch wenn es nicht eigenständig vorkommt. Der Meinung von Liu (2015: 38 f.) nach erscheinen die 17 von Chiao/Kelz genannten Vokalphoneme nicht überzeugend, da viele Allophone mitgezählt wurden. Bezüglich der Aussprache chinesischer Deutschlernender rücken insbesondere die sechs Vokale /a/, /o/, /ɛ/, /i/, /u/, /y/ in den Fokus der Betrachtung (siehe auch Kapitel 2.3.3). Diese bilden eine wichtige Spur, da Abweichungen in ihrer Perzeption und Produktion häufig die Grundlage für systematische orthographische Fehlermuster im Deutschen darstellen.

Im Vergleich zum Deutschen ist die chinesische Sprache reich an Vokalverbindungen (vgl. Hachenberg, 2003). Es gibt im Chinesischen vier fallende Diphthonge: /ai/ (wie <ai> in *bái*), /au/ (wie <ao> in *hǎo*), /ɛi/ (wie <ei> in *fěi*), /ou/ (wie <ou> in *gǒu*). Außerdem bilden die Übergangslaute /i/, /u/, /y/ mit den silbentragenden Vokalen die diphthongischen Einheiten /i̯a/, /i̯ɛ/ (wie <ia> in *jīā*), /i̯ə/ (wie <ie> in *jiě*), /u̯o/ (wie <uo> in *guō*), /u̯a/ (wie <ua> in *guā*) und /y̯ɛ/ (wie <üe>² in *yùè*) (vgl. Hunold, 2021, S. 3). Darüber hinaus werden vier Triphthonge aus den fallenden Diphthongen und den Übergangslauten /i/ und /u/ gebildet: /u̯ai/ (wie <uai> in *shuài*), /i̯au/ (wie <iao> in *biǎo*), /u̯ɛi/ (wie <ui> in *shuǐ*) und /i̯ou/ (wie <iu> in *jǐu*) (vgl. Hunold, 2009, S. 89; T. Liu, 2015, S. 38; Tang, 2003, S. 127). Hier werden bei der Verschriftlichung der beiden Triphthonge /u̯ɛi/

2 Nach *Pinyin*-Orthographie wird Diphthong /y̯ɛ/ nach *j, q, x, y* statt <üe> als <ue> verschriftlicht.

und /jou/ wiederum die *Pinyin*-Orthographieregelungen verwendet. In der folgenden Tabelle werden die Phonem-(*Pinyin*)-Graphem-Beziehungen jeweils im Chinesischen und Deutschen dargestellt.

Tabelle 7: Phonem-Graphem-Korrespondenzen der Vokale und Vokalverbindungen im Deutschen und Chinesischen (*Pinyin*) (Quelle: In Anlehnung an T. Liu, 2015, 38 f.)

Graphemisches Inventar (Deutsch)	Phonematische Basis (IPA)	<i>Pinyin</i> -Graphem	Vorkommen des Vokals in <i>Pinyin</i>
Monophthonge			
<i>, <ie>, <ieh>	/i:/		
	/i/	<i>	allein als Auslaut und in allen Verbindungen, außer nach z-, c-, s-, zh-, ch-, sh-, r-
<i>	/ɪ/		
	/y/	<u> oder <ü>	<u> nach j-, q-, x- und nach dem Übergangsvokal y-; <ü> allein im Auslaut und in allen Verbindungen
<ü>, <üh>	/y:/		
<ü>	/ʏ/		
<e>, <ä>	/ɛ/	<e> oder <a>	<e> im Auslaut nach i-, ü- und vor -i <a> in den Auslauten -ian und -uan nach j-, q-, x-
<e>, <ee>, <eh>	/e:/		
<ä>, <äh>	/ɛ:/		
<ö>, <öh>	/ø:/		
<ö>	/œ/		
<e>	/ə/	<e>	in den Auslauten -en und -eng
	/ə/	<er>	in der Silbe er
<er>	/ɐ/		
<a>	/a/	<a>	allein als Auslaut
<a>, <aa>, <ah>	/a:/		
	/u/	<u>	allein im Auslaut und in allen anderen Verbindungen
<u>, <uu>, <uh>	/u:/		
<u>	/ʊ/		
	/ʊ/	<e>	allein als Auslaut
	/o/	<o>	alleinstehend; allein als Auslaut; in den Auslauten -uo und -ou; in den Auslauten -ao, -iao; -iong, -ong
<o>, <oo>, <oh>	/o:/		
<o>	/ɔ/		

(Fortsetzung Tabelle 7)

Graphemisches Inventar (Deutsch)	Phone-matische Basis (IPA)	Pinyin-Graphem	Vorkommen des Vokals in Pinyin
Diphthonge			
<ei>, <ai>	/aɪ/ bzw. /ai/	<ai>	Fallender Diphthong (z. B. bái).
<au>	/aʊ/ bzw. /au/	<ao>	Fallender Diphthong (z. B. hǎo).
<eu>, <äu>	/ɔʏ/		
	/ɛɪ/	<ei>	Fallender Diphthong (z. B. fēi).
	/ou/	<ou>	Fallender Diphthong (z. B. gǒu).
	/i̯a/	<ia>	Steigende diphthongische Einheit (Übergangslaut + Vokal, z. B. jiā).
	/i̯ə/	<ie>	Steigende diphthongische Einheit (z. B. jiě).
	/u̯o/	<uo>	Steigende diphthongische Einheit (z. B. guǒ).
	/u̯a/	<ua>	Steigende diphthongische Einheit (z. B. guā).
	/y̯ɛ/	<üe>	Steigende diphthongische Einheit (z. B. yuè).
Triphthonge			
	/u̯ai/	<uai>	Gebildet aus Übergangslaut und Diphthong (z. B. shuài).
	/i̯au/	<iao>	Gebildet aus Übergangslaut und Diphthong (z. B. biǎo).
	/u̯ɛi/	<ui>	Bei der Verschriftlichung greifen Pinyin-Orthographieregelungen (Verkürzung, z. B. shuǐ).
	/i̯ou/	<iu>	Bei der Verschriftlichung greifen Pinyin-Orthographieregelungen (Verkürzung, z. B. jiǔ).

Anhand des Vergleiches der Vokalsysteme beider Sprachen lassen sich folgende Unterschiede zusammenfassen (vgl. Huang & Liao, 2017, 101 f.; Hunold, 2009, 87 f., 2021, 4 f.; Qian, 2006, 14 f.; Tang, 2003, S. 126):

- Das Chinesische verfügt über ein deutlich kleineres Inventar an Vokalphonemen, was in Tabelle 7 durch die zahlreichen Leerstellen in der Pinyin-Spalte visuell deutlich wird.
- Im Gegensatz zum Deutschen ist die Vokallänge im Chinesischen nicht phonematisch funktionalisiert, d. h. lange und kurze Vokale sind im Chinesischen nicht distinktiv. Im Vokalsystem des Deutschen hingegen steht jedem langen Vokal ein kurzer gegenüber. Dabei geht die Längenopposition mit einem Qualitätsunterschied (gespannt vs. ungespannt) einher, sodass strukturell das kurze, ungespannte /ɛ/ das phonologische Gegenstück zum langen, gespannten /e:/ bildet.

- Der chinesische ungerundete Hinterzungenvokal /ɤ/ (im Pinyin als <e> verschriftlicht) besitzt im deutschen Vokalsystem generell keine Entsprechung. Wie aus der Tabelle ersichtlich wird, bleibt die Spalte des deutschen Inventars hier leer, da das Deutsche im Bereich der hinteren Vokale ausschließlich gerundete Laute (wie /o/ oder /u/) aufweist.
- Im Chinesischen kommt das gerundete vordere Vokalpaar /ø:/ bzw. /œ/ nicht vor.
- Während das chinesische Phonemsystem ein umfangreicheres Inventar an Diphthongen und Triphthongen aufweist, finden sich im Deutschen keine Triphthonge. Außerdem besitzt der deutsche Diphthong /ɔɪ/ (z. B. in neu) keine direkte Entsprechung im Chinesischen.

Nach IPA lassen sich Konsonanten nach drei Merkmalsdimensionen anordnen: horizontal nach dem Artikulationsort, vertikal nach der Artikulationsart und paarweise nach dem Merkmal der Stimmhaftigkeit (vgl. Fuhrhop & Peters, 2013, S. 29). In der Literatur variiert die Zahl der deutschen Konsonanten zwischen 22 und 25 (vgl. Eisenberg, 2013, S. 85; Fuhrhop & Peters, 2013, S. 32; Hunold, 2021, 5 f.). Die unterschiedliche Besetzung der Konsonanteninventare unterscheidet sich vor allem bei den Lautvarianten von [ʁ] („geriebenes r“), [r] („Zungenspitzen-r“), [ʀ] („Zäpfchen-r“) sowie von [χ] und [x]. Um die Konsonantensysteme des Deutschen vollständig darstellen zu können, wird im folgenden Schema das deutsche Konsonanteninventar mit 25 Konsonanten nach Artikulatoren gezeigt.

Tabelle 8: Konsonanten des Deutschen (Quelle: Fuhrhop & Peters, 2013, S. 32)

	labial		koronal			dorsal			glottal
	bilabial	labio-dental	dental	alveolar	postal-veolar	palatal	velar	uvular	glottal
Plosiv	p b		t d				k g		ʔ
Frikativ		f v		s z	ʃ ʒ	ç ʝ	x	χ ʁ	h
Nasal	m		n				ŋ		
Vibrant			r					R	
Lateraler Approximant			l						

Die paarweise am gleichen Artikulationsort und in der gleichen Artikulationsart in der Tabelle auftretenden Konsonanten werden häufig als *Fortis-Lenis-Paare* bezeichnet. Die Unterscheidung zwischen Fortis und Lenis bezieht sich auf die Artikulationsstärke. Unter den Plosiven gelten /p/, /t/ und /k/ als Fortis-Konsonanten und /b/, /d/, /g/ als Lenis-Konsonanten (vgl. Fuhrhop & Peters, 2013, 32 f.). Im Hochdeutschen ist die Opposition zwischen Fortis und Lenis im Silbenauslaut aufgehoben. Dieses Phänomen wird als *Auslautverhärtung* bezeichnet; darunter versteht man die lautliche Realisie-

rung stimmhafter Obstruenten (Plosive, Affrikaten und Frikative) am Silbenende als deren stimmlose Entsprechungen (vgl. Eisenberg, 2013, S. 6).

Im Chinesischen gibt es 16 Konsonantenphoneme und sechs Affrikaten (vgl. Hunold, 2021, S. 5). Zwar verfügt auch das Deutsche über Affrikaten (wie /pf/ oder /ts/), im chinesischen System nimmt diese Lautklasse jedoch einen deutlich breiteren Raum ein. Es gibt in der chinesischen Sprache keine Fortis-Lenis-Paare. Entscheidend und bedeutungsunterscheidend ist hier das Merkmal der Aspiration. So stehen den stimmlosen, nicht aspirierten Explosiven /p/, /t/, /k/ die aspirierten Explosive /p^h/, /t^h/ und /k^h/ gegenüber. Bei den Frikativen fehlt jedoch diese Opposition (vgl. Hunold, 2009, S. 89, 2021, S. 5; Qian, 2006, S. 14). Auf der graphematischen Ebene werden diese Phoneme im *Hanyu Pinyin Fang'an* durch 21 Konsonanten-(*Pinyin*-)grapheme repräsentiert: , <p>, <m>, <f>, <d>, <t>, <n>, <l>, <g>, <k>, <h>, <j>, <q>, <x>, <zh>, <ch>, <sh>, <r>, <z>, <c>, <s>.

Tabelle 9: Konsonanten des Chinesischen (Quelle: In Anlehnung an Hunold 2009, S. 89; Hunold 2021, S. 5)

	bilabial	labio-dental	alveolar	retro-flex	palatal	velar
stimmlose/nicht aspirierte Explosive	p		t			k
stimmlose/aspirierte Explosive	p ^h		t ^h			k ^h
stimmlose Frikative		f	s	ʂ	ɕ	x
stimmhafte Frikative				ʐ		
stimmlose/nicht aspirierte Affrikate			ts	tʂ	tɕ	
stimmlose/aspirierte Affrikate			ts ^h	tʂ ^h	tɕ ^h	
Nasale	m		n			ŋ
Liquide			l			

Die Gegenüberstellung der Konsonanten-(*Pinyin*-)Graphem-Beziehungen in Tabelle 10 verdeutlicht die Korrespondenzen zwischen dem deutschen Schriftsystem und dem *Pinyin*-System. Zur Veranschaulichung der praktischen Anwendung der Konsonanten werden dabei illustrative Wortbeispiele sowie deren entsprechende lautliche Umschrift (in Klammern) angeführt. Angesichts der Meinung von Hunold (2009, S. 89), dass sich die sechs Affrikaten im Chinesischen phonologisch als Folge zweier Phoneme betrachten lassen, werden die vier Konsonantenverbindungen im Deutschen /pf/, /ts/, /tʃ/, /dʒ/ auch in der Tabelle aufgelistet.

Tabelle 10: Konsonanten-(*Pinyin*)-Graphem-Beziehungen im Chinesischen und im Deutschen (Quelle: In Anlehnung an T. Liu, 2015, 45 f.)

Deutsche Konsonanten-graphem(e)	Konsonanten (IPA)	<i>Pinyin</i> -Graphem	Vorkommen des Konsonanten im <i>Pinyin</i> -Wort
<p>, <pp>, 	/p/		bù (/pù/)
	/p ^h /	<p>	pà (/p ^h a/)
, <bb>	/b/		
<m>, <mm>	/m/	<m>	mā (/ma/)
<ff>, <f>, <v>, <ph>	/f/	<f>	fā (/fa/)
<w>, <v>	/v/		
<pf>	/pf/		
<t>, <tt>, <th>, <d>, <dd>, <dt>	/t/	<d>	dé (/tr/)
	/t ^h /	<t>	tā (/t ^h a/)
<d>	/d/		
<s>, <ß>, <ss>	/s/	<s>	sì (/sɿ/)
<s>	/z/		
<z>, <tz>, <ts>, <c>	/ts/	<z>	zǐ (/tsɿ/)
	/ts ^h /	<ci>	cǐ (/tɕ ^h ɿ/)
<r>, <rr>	/r/		
<n>, <nn>	/n/	<n>	nà (/na/)
<l>, <ll>	/l/	<l>	lā (/la/)
<sch>, <s>	/ʃ/		
<g>	/ʒ/		
<tsch>	/tʃ/		
<dsch>, <dg>	/dʒ/		
	/ʂ/	<sh>	shì (/ʂɿ/)
	/ʐ/	<r>	rì (/ʐɿ/)
	/tʂ/	<zh>	zhī (/tʂɿ/)
	/tʂ ^h /	<ch>	chī (/tʂ ^h ɿ/)
	/ɕ/	<x>	xiāo (/ɕi̯au/)
<ch>, <g>	/ç/		

(Fortsetzung Tabelle 10)

Deutsche Konsonanten-graphem(e)	Konsonanten (IPA)	Pinyin-Graphem	Vorkommen des Konsonanten im Pinyin-Wort
<j>	/j/		
	/tɕ/	<j>	jiā (/tɕja/)
	/tɕʰ/	<q>	qiāo (/tɕʰja/)
<ck>, <kk>, <k>, <g>	/k/	<g>	gē (/kɤ/)
	/kʰ/	<k>	kè (/kʰɤ/)
<g> <gg>	/g/		
<ch>	/x/	<h>	hào (/xau/)
<ng>	/ŋ/	<ng>	nóng (/nʊŋ/)
<r>	/ʁ/		
<r>	/R/		
<h>	/h/		

Beim Vergleich des Konsonantensystems beider Sprachen sind die folgenden Unterschiede festzustellen (vgl. Hunold, 2009, 90 f., 2021, S. 5; T. Liu, 2015, 48 f.; Tang, 2003, 125 f.):

- Das Konsonanteninventar im Chinesischen ist weniger umfangreich als im Deutschen.
- Im Deutschen stehen die stimmhaften (lenis), nicht aspirierten Konsonanten /b/, /d/, /g/, /v/, /z/, /ʒ/, /j/, /r/ den stimmlosen (fortis) Konsonanten /p/, /t/, /k/, /f/, /s/, /ʃ/, /ç/ und /x/ gegenüber. Im Chinesischen hingegen existiert diese Unterscheidung zwischen Fortis und Lenis nicht. Stattdessen unterscheidet das Chinesische zwischen schwach gespannten, nicht aspirierten Konsonanten (/p/, /t/, /k/, /tɕ/, /tɕʰ/, /ts/) und stark gespannten, stark aspirierten (/pʰ/, /tʰ/, /kʰ/, /tɕʰ/, /tɕʰʰ/, /tsʰ/). Dabei spielt der Grad der Aspiration eine zentrale Rolle.
- Der retroflexe Konsonant /ɣ/ im Chinesischen und das deutsche /ʃ/ ähneln sich zwar, sind jedoch nicht identisch; /ɣ/ ist ein retroflexes Phonem.
- Im Chinesischen gibt es nur den stimmlosen velaren Frikativ /x/, während das Deutsche den stimmlosen glottalen Frikativ /h/ verwendet, obwohl beide durch das Graphem <h> repräsentiert werden.
- Der stimmhafte labiodentale Frikativ /v/ existiert im Chinesischen nicht, dort gibt es nur /f/; die Frikative /z/, /ʒ/, /j/ kommen im Chinesischen ebenfalls nicht vor.
- Von den Affrikaten sind im Chinesischen /pf/, /tʃ/, /dʒ/ unbekannt, jedoch ist /ts/ vorhanden.
- Im Deutschen wird das Graphem <s> sowohl für das stimmlose Phonem /s/ als auch für das stimmhafte Phonem /z/ verwendet. Im Chinesischen hingegen steht das Pinyin-Graphem <s> ausschließlich für das stimmlose Phonem /s/.

- Das konsonantische <r> im Deutschen (stimmhafter uvularer Frikativ [ʁ] und die Vibranten [r], [ʀ]) fehlt im Chinesischen; das chinesische <r> ist ein stimmhaftes retroflexes /ʐ/.
- Retroflexe Konsonanten wie /ʂ/, /ʐ/, /tʂ/, /tʂʰ/ sowie alveolopalatale Konsonanten wie /tɕ/ und /tɕʰ/ kommen im Deutschen nicht vor.

Abgesehen von den Unterschieden bei den Phoneminventaren weicht die Phonotaktik der beiden Sprachen sehr stark voneinander ab. Als Teildisziplin der Phonologie beschäftigt sich die Phonotaktik mit den syntagmatischen Beziehungen zwischen Lauten innerhalb der Silbe und anderer prosodischer Einheiten (vgl. Fuhrhop & Peters, 2013, S. 88). Die phonologische Silbe ist die kleinste Lautfolge, die selbstständig geäußert werden kann. In der Literatur werden vor allem zwei Strukturmodelle für die Silbenanalyse vorgestellt: Konstituentenmodelle und CV-Modelle. Die Konstituentenmodelle identifizieren Strukturpositionen, die in einem hierarchischen Verhältnis zueinander stehen; dabei werden die Silben in Anfangsrand, Silbenkern, Endrand eingeteilt. Die CV-Modelle schreiben der Silbe eine „flache“ Silbenstruktur zu, und Silben werden als Folgen von C- und V-Positionen repräsentiert. V-Elemente stehen für Positionen in der Silbe, die durch silbische Laute besetzt werden; C-Elemente stehen für Positionen, die durch nicht silbische Laute besetzt werden (vgl. Fuhrhop & Peters, 2013, 79 ff.). In dieser Arbeit werden die Konstituentenmodelle für die Fehleranalyse in Bezug auf Silbenstruktur verwendet.

Wie am Anfang von Kapitel 2.3 bereits erwähnt, besteht eine chinesische Silbe aus einem Anlaut am Silbenanfang und einem Auslaut am Silbenende. Qiu (1984b, S. 85) erfand ein Silbenmodell für die Strukturformel der chinesischen Phonotagmen:

$$(K_1)(V_3)V_1\left(\frac{V_2}{K_2}\right)$$

Abbildung 3: Die phonotaktische Struktur der chinesischen Silbe (Quelle: Qiu, 1984b, S. 85)

Der Anlaut kann aus einem Konsonanten (K_1) bestehen, kann aber auch leer sein; in diesem Fall wird er als *Null-Initial* bezeichnet. Den obligatorischen Kern des Auslautes bildet stets ein Hauptvokal – entweder als Monophthong (V_1) oder als Diphthong (V_1+V_2). Vor dem Hauptvokal im Auslaut kann ein weiterer Übergangshalbvokal aus /i/, /u/, /y/ als V_3 auftreten. In der Formel kennzeichnen Klammern optionale Elemente: Während V_1 als obligatorischer Silbenkern ungeklammert erscheint, sind K_1 , V_3 sowie die abschließende Position (V_2/K_2) fakultativ. Die vertikale Notation von V_2 und K_2 zeigt an, dass es sich um zwei sich gegenseitig ausschließende Alternativen an derselben Silbenposition handelt: Nach dem Hauptvokal kann entweder ein zweiter Vokal V_2 auftreten – wodurch ein Diphthong entsteht – oder ein nasaler Konsonant (K_2 /n/, /ŋ/) als Silbenende, nicht jedoch beides gleichzeitig. Silben im Chinesischen können ohne Anlaut realisiert werden, wie z. B. 啊 á /a/ (eine modale Partikel); keine Silbe kann ohne

Auslaut sein. Außerdem kommen Konsonantenhäufungen nicht vor (vgl. Hunold 2009: 72).

Im Vergleich zu der Silbenstruktur im Chinesischen ist das deutsche Silbenmodell komplizierter. Einerseits weist das deutsche Silbenmodell eine weitaus größere Vielzahl an Silbentypen auf, andererseits gibt es innerhalb der einzelnen Silbe auch mehr Möglichkeiten. Pompino-Marschall (2003, 273 f.) bildet nach Kohler (1995) ein Modell für einsilbige Wörter im Deutschen ab:

$$\left(\left(\begin{array}{ccc} & K_{a,b,c} & \\ (K_a) & K_a & K_b \\ & K_a & K_c \\ (K_a) & K_a & K_a \end{array} \right) V \left(\begin{array}{ccc} & K_{a,b} & \\ K_b & K_a & (K_a) \\ K_b & K_b & (K_a) \\ K_a & K_a & \end{array} \right) \left(\left(\begin{array}{cc} K_a & (+K_a) \\ (+K_a) & (K_a) \end{array} \right) \right)$$

Abbildung 4: Die phonotaktische Struktur des deutschen Einsilbers (Quelle: Pompino-Marschall, 2003, 273 f.)

K_a bezeichnet plosive oder frikative Konsonanten; K_b bezeichnet Nasale oder /l/ oder /ʁ/; K_c beschreibt die Konsonanten /h/ und /j/, und V sind monophthongische und diphthongische Silbenkernsegmente. Die runden Klammern in der Abbildung kennzeichnen dabei fakultativ auftretende Segmente, während in geschweiften Klammern Wahlmöglichkeiten der untereinander angeordneten Segmente dargestellt sind. Das Zeichen „+“ gibt eine Morphemgrenze (wie z. B. in *schimpf+st*) an. Hiernach ist das Vokalsegment in der Mitte der einzig nötige Bestandteil einer deutschen Silbe. Im Vergleich zum Chinesischen sind die Konsonantenhäufungen in verschiedenen Positionen einer deutschen Silbe möglich: Es können am Wortbeginn bis zu drei Konsonanten auftreten, wie u. a. /str/ in *Strick*, /pfr/ in *Pfropf*, /pfl/ in *Pflock*. Unter Einschluss einer Morphemgrenze können am Wortende bis zu fünf Konsonanten aufeinandertreffen, wie z. B. /rbsts/ in (*des*) *Herbsts*. Für die finalen Konsonanten der Klasse K_a besteht auf Grund der deutschen Auslautverhärtung die Einschränkung auf stimmlose Konsonanten (vgl. Pompino-Marschall, 2003, S. 273).

Benedix (2009, S. 90) fasst die phonotaktischen Möglichkeiten der beiden Sprachen zusammen und stellt fest, dass sich aus der chinesischen Silbenstruktur nur 10 Möglichkeiten ergeben, während die deutsche Silbenstruktur über 35 Möglichkeiten verfügt. In vielen Arbeiten wird die Komplexität der deutschen Phonotaktik als Schwierigkeit beim Fremdspracherwerb der chinesischen Deutschlernenden erklärt (siehe nächstes Kapitel).

Aus der Perspektive der CDST werden Ausspracheabweichungen nicht als einfache Interferenzen aus L1 verstanden, sondern als Ergebnis einer komplexen Wechselwirkung zwischen den bereits vorhandenen sprachlichen Ressourcen und der Zielsprache Deutsch. Bislang existieren jedoch kaum Untersuchungen, die Transferphänomene bei chinesischen DaF-Lernenden direkt auf der orthographischen Ebene analysieren (siehe Kapitel 2.4). Da die deutsche Orthographie in weiten Teilen auf dem phonographischen Prinzip beruht, wirken sich die phonetisch-phonologischen Interaktionen direkt auf die Rechtschreibleistung der Lernenden aus. Die bestehenden Forschungen zu Ausspracheabweichungen können daher wertvolle Hinweise auf potenzielle Transferphänomene

auf orthographischer Ebene liefern. Das folgende Kapitel stellt aus diesem Grund die wichtigsten in der Forschungsliteratur dokumentierten Aussprachephänomene bei chinesischen Deutschlernenden vor.

2.3.3 Potenzielle Transferphänomene zwischen dem Chinesischen und Deutschen auf orthographischer Ebene

In den meisten Untersuchungen (siehe Kapitel 2.4) wird übereinstimmend festgestellt, dass orthographische Fehler bei chinesischen Deutschlernenden eng mit Schwierigkeiten in der Aussprache verknüpft sind. Hunold (2021, S. 8) identifiziert drei zentrale Lehr- und Lernschwerpunkte in Deutsch als Fremd- und Zweitsprache für Lernende mit Chinesisch als L1:

- Schwerpunkt 1: Suprasegmentalia
- Schwerpunkt 2: Segmentalia: Vokale & Konsonanten
- Schwerpunkt 3: Phonem-Graphem-Beziehungen

Im Folgenden wird die bestehende Forschung zu Ausspracheproblemen chinesischer Deutschlernender analysiert. Das Ziel ist es, herauszufinden, welche der genannten Schwerpunkte als Ausdruck dynamischer Anpassungsprozesse und des L1-Einflusses identifiziert werden können. Auf diese Weise soll die enge Korrelation zwischen Ausspracheabweichungen und orthographischen Fehlern verdeutlicht werden.

Die suprasegmentale Ebene umfasst Aspekte wie Akzent, Rhythmus und Melodie. Genau hier zeigen sich bei chinesischen Deutschlernenden Schwierigkeiten, die sich aus dem fundamentalen typologischen Kontrast zwischen ihrer L1 und der Zielsprache ergeben. Das Deutsche funktioniert als Akzentsprache, die Informationen durch die Betonung von Silben in Wörtern und Sätzen strukturiert. Im Gegensatz dazu ist das Chinesische eine Tonsprache, in der die melodische Kontur einer Silbe deren Bedeutung bestimmt (Hunold, 2021, 3 ff.). In ihrer Untersuchung zu suprasegmentalen Abweichungen der Aussprache chinesischer Deutschlernender belegt Hunold (2009) eine unbewusste Übertragung des phonologischen Musters vom Chinesischen auf das Deutsche. Dies führt dazu, dass das Verbinden mehrerer Silben zu einem Wort oder mehrerer Wörter zu Sprechereinheiten nicht im angemessenen rhythmischen und melodischen Zusammenhang erfolgt. Äußerungen werden durch zu viele oder falsche Pausen zergliedert (vgl. Hunold, 2009, S. 155–166). Auch Lay (2017, S. 33) beobachtet dieses Phänomen und nennt als typisches Beispiel das silbenweise Aussprechen wie *Lo-ko-mo-ti-ve* anstatt des gebundenen Wortes *Lokomotive*.

Diese primär suprasegmentalen Abweichungen scheinen auf den ersten Blick nur von geringer Bedeutung für die Orthographie zu sein. Der Spracherwerb ist jedoch ein komplexes dynamisches System, in dem alle Ebenen miteinander interagieren. Eine falsche Akzentuierung kann daher direkte Folgen für die Getrennt- und Zusammenschreibung haben. Wird beispielsweise ein Kompositum wie ['ʃpʁa:xʃu:lə] (*Sprachschule*) fälschlicherweise auf der zweiten Silbe betont [ʃpʁa:x'ʃu:lə] (*Sprachschule*), neigen Lernende dazu, es als zwei getrennte Wörter aufzufassen und zu schreiben (*Sprache schule*) (vgl. Peperkamp & Dupoux, 2002). In ähnlicher Weise werden unbetonte deut-

sche Präfixe (z. B. {ver-} in *verstehen*) unter dem Einfluss der L1 oft fälschlicherweise betont, was dazu führt, dass sie als eigenständige Einheiten wahrgenommen und vom Verbstamm getrennt geschrieben werden (vgl. Bassetti, 2017).

Im Bereich der Phonotaktik nimmt Qiu (1984a) eine ausführliche kontrastive Analyse zwischen den Phonemsystemen des Deutschen und des Chinesischen vor. Dabei liegt die Annahme nahe, dass die phonotaktische Schwierigkeit bei den chinesischen Deutschanfängern häufig durch interlinguale Interferenz bedingt ist (vgl. Qiu, 1984a; Peperkamp & Dupoux, 2002). Ein markanter Unterschied besteht darin, dass im Chinesischen keine Konsonantencluster vorkommen. Dies führt zu vielen Schwierigkeiten, insbesondere bei der Aussprache von Endplosiven, wo häufig ein zusätzliches [ə] eingefügt wird, z. B. [niçtə] statt [niçt] (vgl. Qiu 1984a). Solche Ausspracheprobleme können sich auch in der Schrift manifestieren (vgl. Qiu, 1984a, S. 126). Auch Hunold weist ebenfalls darauf hin, dass Konsonantenhäufungen eine große Herausforderung für chinesische Deutschlernende darstellen und eine häufige Fehlerquelle sind. (Vgl. Hunold, 2009, S. 168)

M. Wang (1993) betrachtet die segmentalen Ausspracheprobleme chinesischer Deutschlernender umfassend und unterstreicht ebenfalls den Einfluss der phonotaktischen Besonderheiten des Chinesischen. Ein häufiges Phänomen ist das silbenweise Aussprechen deutscher Wörter oder Sätze, was Wang auf die traditionelle chinesische Phonologie zurückführt. In dieser wird der anlautende Konsonant einer chinesischen Silbe als *Initial* und der folgende Vokal oder Nichtkonsonant als *Final* klassifiziert (vgl. M. Wang, 1993, 66 ff.). Im Deutschen hingegen ist die Silbenstruktur nicht immer mit einem Phonotagma identisch. Diese Ausspracheproblematik könnte zu Fehlern bei der Getrennt- und Zusammenschreibung deutscher Wörter führen.

Zu den segmentalen Abweichungen chinesischer Deutschlernender finden sich in der Literatur ausführlichere Auseinandersetzungen. Häufig wird das deutsche Vokalsystem als ein besonders problematisches für chinesische Lernende beschrieben. Dabei werden die Vokallänge und -kurze als zentrale Fehlerquellen bezeichnet (vgl. Hachenberg, 2003; Hunold, 2009, 2021, S. 8; M. Wang, 1993; Yen, 1992). Für Hachenberg (2003) stellen die Lang- und Kurzvokale der deutschen Sprache sowohl in der Wahrnehmung als auch in der Produktion eine zentrale Schwierigkeit dar. Daneben nennt er die Verwechslung der beiden Schwa-Laute [ɐ] und [ə] als häufige Fehler. Ebenso treten Verwechslungen zwischen den langen und kurzen Phonemen /y:/ und /ʏ/ sowie /ø:/ und /œ/ auf, als auch zwischen den langen Vokalen /o:/ und /u:/. Besonders auffällig ist die diphthongierte Aussprache von /ɛ/ zu [eɪ], die ebenfalls zu Rechtschreibfehlern führen kann (vgl. Hachenberg, 2003, 103 ff.). Hunold (2009, 166 f.; Hunold, 2021, S. 8) beobachtet zudem, dass chinesische Lernende häufig kurze Vokale zu lang aussprechen, was bei allen kurzen Vokalen auffällt. Umgekehrt sind bei langen Vokalen die Qualitätsabweichungen bemerkenswert, insbesondere beim Vokal /o:/, der offen als [ɔ] realisiert wird. Auch beim Vokal /e:/ zeigen sich qualitative Abweichungen, die entweder in eine offene Aussprache als [ɛ] oder in eine diphthongierte Form als [eɪ] münden.

M. Wang (1993) hingegen sieht die Vokale als weniger problematisch und betont, dass die hauptsächlichen Schwierigkeiten in der Aussprache der deutschen Konsonan-

ten liegen. Sprenger (1974, 183 f.) hebt hervor, dass insbesondere die im Chinesischen fehlenden Konsonanten /r/, /h/ und /v/ problematisch sind. Oftmals kommt es bei chinesischen Lernenden zu einer Nasalisierung des vorangehenden Vokals, wie z. B. für *München*: [ˈmʏnɐən] wird [ˈmʏɐən] (*<Müchen>) gesprochen. Dies kann auf der schriftlichen Ebene zur Weglassung des Konsonantengraphems <n> führen.

Die Schwierigkeit mit dem R-Laut wird auch von C. Liu (1982), Spiekermann (1997) und Hachenberg (2003) bestätigt. Hier kommt es häufig zu einer Substitution von /r/ durch [l], oder das /r/ wird durch ein stark reibendes [x] oder durch [ʒ] (*Pinyin*: <r>) ersetzt. Auf der schriftlichen Ebene kann dies zu Verwechslungen bei der Graphem-Phonem-Korrespondenz führen, wie zwischen <r> und <l> oder <r> und <h>.

Hachenberg macht erstmals auf die Interferenz des *Pinyin*-Systems bei der Analyse der Lernersprache aufmerksam. So wird der deutsche Konsonant /ʃ/ oft durch das *Pinyin*-Phonem /ɕ/ (<sh>) ersetzt. Auch das Fehlen des Lautes [ç] im Chinesischen führt häufig dazu, dass dieser durch das alveolo-palatale /ɕ/ (<x>) substituiert wird (vgl. Hachenberg, 2003, 119 f.).

Hinsichtlich der genauen Kausalität dieser Fehlerentstehung lässt sich methodisch oft nur schwer trennen, ob eine primär abweichende Lautwahrnehmung zur falschen Verschriftlichung führt oder ob verinnerlichte Grapheme des *Pinyin*-Systems die Artikulation steuern. Vielmehr ist in solchen Fällen von einer reziproken phonologisch-graphematischen Überlagerung auszugehen. Wenn Lernende das deutsche Wort *Schule* /ʃu:lə/ durch den phonologischen Filter ihrer Erstsprache wahrnehmen, greifen sie bei orthographischer Unsicherheit offenbar auf die ihnen hochgradig automatisierte Graphem-Phonem-Korrespondenz des *Pinyin* (hier <sh> für den Laut [ʃ]) zurück. Schreibungen wie *<shule> statt des korrekten <Schule> lassen sich in der Fehleranalyse folglich als emergentes Resultat dieser L1-basierten Interferenz interpretieren. Die aus der L1 übertragene (*Pinyin*-)Graphem-Phonem-Korrespondenz fungiert somit als eine wesentliche kognitive Quelle für diese orthographischen Abweichungen.

Die beschriebenen Ausspracheschwierigkeiten und ihre möglichen Auswirkungen auf die orthographische Leistung chinesischer Deutschlernender deuten darauf hin, dass die im multikompetenten System der Lernenden verankerte Erstsprache Chinesisch (inklusive ihrer phonetischen Repräsentation durch *Pinyin*) als einflussreicher Faktor die Entwicklung der deutschsprachigen orthographischen Kompetenz mitgestaltet. Im nachfolgenden Kapitel 2.4 wird der allgemeine Forschungsstand zur Rechtschreibfehleranalyse im DaZ/DaF-Kontext dargestellt.

2.4 Forschungsstand zur Rechtschreibfehleranalyse im DaZ/DaF-Kontext

„Verglichen mit anderen Aspekten des Deutschen spielt die Rechtschreibung in den Lehrmaterialien Deutsch als Fremdsprache keine oder nur eine sehr untergeordnete Rolle“ (Rösler, 1982, S. 29). Diese bereits vor über vier Jahrzehnten getroffene Feststellung wurde seither wiederholt von Expert*innen im Bereich Deutsch als Fremdsprache

bestätigt (vgl. Eisenberg, 1995:171; Földes, 2000; Hans-Bianchi, 2007; Hans-Bianchi & Katelhön, 2010; Hayes-Harb, Brown & Smith, 2018) und hat auch heute noch weitgehend Gültigkeit. Außerdem bleibt der Orthographieerwerb von DaZ/DaF-Lernenden auch im wissenschaftlichen Diskurs ein Randthema. Die wenigen existierenden Forschungsarbeiten hierzu fokussieren überwiegend auf Lernende mit europäischen Erstsprachen wie Türkisch (Becker, 2013; 2018; Nimz, 2022; Thomé, 1987), Spanisch (Corvacho Del Toro, 2004), Russisch (H.-G. Müller & Schroeder, 2022) oder Italienisch (Hans-Bianchi, 2007; Hans-Bianchi & Katelhön, 2010; Richter, 2008).

Seit den 1980er-Jahren erscheinen deutlich mehr Studien zum Orthographieerwerb bilingual aufwachsender Schüler*innen. Diese Untersuchungen konzentrieren sich vor allem auf Rechtschreibfehler und analysieren zentrale orthographische Phänomene in den Lernertexten. Viele Studien stellen zudem Vergleiche der Rechtschreibkompetenz mehrsprachiger Schüler*innen mit jener ihrer deutsch-monolingualen Mitschüler*innen an.

Thomé (1987) analysiert 3.450 Rechtschreibfehler aus frei formulierten Klassenarbeiten von 61 Schülerinnen mit Türkisch als L1 und 28 deutschen Schülerinnen der 8. Klasse aus zwei Berliner Gesamtschulen. Er unterteilt die orthographischen Fehler in 40 Einzeltypen. Die prozentuale Darstellung der Fehlerhäufigkeiten zeigt, dass die türkischen Schüler*innen die meisten Schwierigkeiten im Bereich der Groß- und Kleinschreibung (36,32 %) aufweisen. Darauf folgen Fehler bei der Markierung von kurzen und langen Vokalen (13,29 %). Bemerkenswert war, dass die sechs häufigsten Fehler der türkischen Schüler*innen (Klein- für Großschreibung; Groß- für Kleinschreibung; Getrennt- für Zusammenschreibung; Konsonant weglassen; s für ß; Vokalkürze nicht markiert) auch die sechs häufigsten Fehler der deutschen Schüler*innen darstellten. Thomé (1987) führt einige Fehler auf Berliner dialektale Einflüsse zurück und stellt fest, dass sich nur wenige Fehler der türkischen Schüler*innen mit der Kontrastivhypothese erklären ließen. Er schlussfolgert zudem, dass bei Kindern, die bereits in ihrer Erstsprache alphabetisiert wurden, graphematische Interferenzen häufiger zu erwarten seien als bei jenen, die ausschließlich Bildungseinrichtungen im Zielland besucht haben.

Zu ähnlichen Ergebnissen hinsichtlich eines begrenzten L1-Einflusses kam Grießhaber (2004) in seiner Untersuchung zum Schriftspracherwerb von Grundschüler*innen an Frankfurter Schulen. An der Studie nehmen 175 Schüler*innen mit unterschiedlichen Erstsprachen teil (Türkisch, Serbokroatisch, Rumänisch, Slawisch und weitere Sprachen). Bei der Kategorisierung der Schreibproben nach Herkunftssprache stellen sich leichte Unterschiede heraus: Kinder, deren Erstsprache Serbokroatisch war, weisen eine stärkere Lautorientierung in ihrer Schreibung auf als Kinder, deren Erstsprachen Deutsch oder Türkisch sind. Dies deutet auf einen positiven Einfluss der Erstsprache hin, die den Kindern mit Serbokroatisch als Erstsprache „ein gutes Wahrnehmungsraster zur lautlichen Erschließung des Deutschen“ bietet (vgl. Grießhaber, 2004, 83 f.). Insgesamt sind direkte Einflüsse der erstsprachlichen Schriftsprachkonventionen zwar im Einzelfall erkennbar, spielen jedoch für die Gesamtleistung eine untergeordnete Rolle (vgl. Grießhaber, 2004, S. 84).

Auch Jeuk (2012; Jeuk, 2021), der 63 Texte von zehn Erstklässler*innen mit Türkisch als Erstsprache analysiert, kommt zu dem Schluss, dass keine signifikanten Unterschiede in den orthographischen Fehlern zwischen ein- und mehrsprachigen Schüler*innen feststellbar seien (vgl. Jeuk, 2012, S. 117).

Im Kontrast zu diesen Befunden stehen jedoch zahlreiche Studien, die deutliche Unterschiede in der orthographischen Kompetenz zwischen DaZ-Lernenden und Schüler*innen mit Deutsch als Erstsprache belegen. Eine häufig zitierte Studie hierzu stammt von Fix (2002). Fix analysiert Rechtschreibfehler in Aufsätzen von 356 Schüler*innen der 8. Klasse (Haupt- und Realschule) und klassifiziert diese Fehler in drei Kategorien: Kategorie A umfasste grammatisch bedingte Fehler, wie Groß- und Kleinschreibung, Getrennt- und Zusammenschreibung sowie die Schreibung von *das/dass* und Wortfamilien mit Auslautverhärtung, <ä> – <e> und <äu> – <eu>. In Kategorie B ordnet Fix Fehler bei der Markierung der Vokalquantität, der s-Schreibung sowie weitere Probleme auf der Laut-Buchstaben-Ebene ein. Kategorie C beinhaltete Fehler bei der Fremdwortschreibung sowie sonstige Rechtschreibfehler (Fix, 2002, 49 f.). Nach der Fehlerauswertung stellt Fix fest: „zwischen den Schularten zeigt sich einen deutlichen Unterschied in den Rechtschreibleistungen. Hauptschüler machen bei wesentlich geringerer Textlänge mehr Fehler“ (Fix, 2002, S. 48). Zudem nimmt Fix an, dass Schüler*innen mit Deutsch als Zweitsprache vor allem Schwierigkeiten im grammatisch-morphologischen Bereich, insbesondere bei der Groß- und Kleinschreibung, hätten. Im phonologischen Bereich seien die Fehler jedoch bei den muttersprachlichen deutschen Schüler*innen sogar stärker ausgeprägt als bei den DaZ-Lernenden.

Ein wiederkehrendes Thema, das auf Leistungsunterschiede hinweist, ist die Vokalquantität. Steinig, Betzel, Geider und Herbold (2009) untersuchen die Schreibkompetenzen von Viertklässlern im diachronen Vergleich (1972 vs. 2002) und finden heraus, dass die Markierung langer und kurzer Vokale den größten Unterschied zwischen bilingualen und monolingualen Kindern ausmacht. Diese Schwierigkeit bilingualer Schüler*innen bestätigt Becker (2013) in ihrer qualitativen Längsschnittstudie (Klasse 1–4) mit neun türkischsprachigen und elf monolingual deutschen Schüler*innen. Die Ergebnisse zeigen signifikante Unterschiede in der Markierung der Vokalgespanntheit: die Schüler*innen mit Türkisch als L1 zeigen große Schwierigkeiten bei der Markierung von Vokallänge und Vokalkürze, und diese Herausforderung besteht während ihrer gesamten Schulzeit. Auch H.-G. Müller und Schroeder (2022) kommen in ihrer Studie zu ähnlichen Ergebnissen bezüglich der Verschriftlichung von langen und kurzen Vokalen. Sie untersuchen den Orthographieerwerb von Schüler*innen im Alter von acht bis zehn Jahren mit Russisch oder Türkisch als Erstsprache. Die Studie basiert auf einem Korpus von orthographischen Übungsdaten von etwa 800 Testschüler*innen, die über die Website *Orthografietrainer.net*³ gesammelt werden. Die Übungen umfassen verschiedene orthographische Herausforderungen, darunter die Markierung von Vokallänge, sowie Aufgaben zur Groß- und Kleinschreibung, zur Getrennt- und Zusammenschreibung usw. Ein zentrales Ergebnis der Studie besteht darin, dass DaZ-Schü-

3 Das Onlineportal *Orthografietrainer.net* (<https://orthografietrainer.net/index.php>) wurde im Rahmen des Forschungsprojekts *OrthographieErwerb* entwickelt, um den Rechtschreiberwerb von Schüler*innen zu erforschen. (H. G. Müller (2008)).

ler*innen tendenziell mehr Rechtschreibfehler machen als ihre Mitschüler*innen, für die Deutsch die L1 ist. Besonders auffällig ist, dass sich diese Unterschiede vor allem bei Aufgaben zur Unterscheidung von gespannter und ungespannter Vokallänge zeigen. In anderen Bereichen der Rechtschreibung sind die Unterschiede weniger stark ausgeprägt. Müller und Schroeder schlussfolgern, dass eine Erstsprache, die keine systematische Unterscheidung zwischen gespannten und ungespannten Vokalen kennt, wie beispielsweise Russisch oder Türkisch, zu entsprechenden Schwierigkeiten bei der Diskriminierung führt. Dies verursacht wiederum spezifische Probleme beim Schreiben (vgl. ebd.: 70).

Die Schwierigkeiten bei der Markierung der Vokalquantität zeigen sich nicht nur bei Schüler*innen mit Türkisch oder Russisch als Erstsprache, sondern auch bei Lernenden mit Italienisch als L1. Richter (2008) führt eine Längsschnittstudie mit deutsch-italienischen Schüler*innen der 4., 5., 6. und 9. Klasse an einer deutsch-italienisch bilingualen Schule durch. Die Untersuchung ergibt, dass die zweisprachigen Schüler*innen zwar qualitativ ähnliche Entwicklungsprozesse wie die einsprachige Vergleichsgruppe durchlaufen, dabei jedoch eine deutliche Verzögerung aufweisen. Besonders in der Repräsentation von gespannten Vokalen, wie der Schreibung mit Dehnungs-*<h>* und der Doppelvokalschreibung, liegen die mehrsprachigen Schüler*innen laut Richter um zwei bis drei Jahre hinter ihren monolingualen Mitschüler*innen zurück.

In den letzten Jahren haben Forschungen zum Orthographieerwerb im DaZ-Kontext verstärkt auf korpuslinguistische Methoden zurückgegriffen. Hierbei werden die Lernertexte transkribiert, als computerlesbares Korpus erstellt, lernersprachliche Fehler annotiert und im Anschluss statistisch analysiert. Ein repräsentatives Beispiel für diesen methodischen Ansatz stellt die Studie von Nimz (2022) dar, die eine Sekundäranalyse des MULTILIT-Korpus (Schroeder & Schellhardt, 2015) durchführt. Ziel der Untersuchung ist die Analyse orthographischer Fehler von Schülerinnen der 5., 7. und 10. Klasse an unterschiedlichen Schulen in Berlin. Dabei werden die Fehler von einsprachig deutsch aufgewachsenen ($n = 28$) und deutsch-türkisch bilingualen ($n = 55$) Schülerinnen verglichen. Die Annotation der orthographischen Fehler erfolgt in EXMARALDA (T. Schmidt & Wörner, 2014)⁴ und umfasst sieben Kategorien, darunter phonographische Abweichungen (PHO), silbische Abweichungen (SIL), morphologische Abweichungen (MOR), syntaktische Abweichungen (SYN) sowie Abweichungen in der Zeichensetzung (ZEI). Die statistische Auswertung ergibt, dass die bilingualen Schüler*innen durchschnittlich mehr Fehler machen als die monolingualen. Dieser Unterschied ist jedoch unter Berücksichtigung von Kontrollfaktoren wie Geschlecht, Testsorte und Klassenstufe statistisch nicht signifikant. Signifikante Unterschiede treten bei syntaktischen Schreibungen auf, insbesondere bei der Groß- und Kleinschreibung sowie der Getrennt- und Zusammenschreibung.

Die Studie von Ruppert und Hanulíková (2022) untersucht ebenfalls Rechtschreibleistungen und Einflussfaktoren anhand von Satzdictaten von 130 Jugendlichen (Deutsch als Muttersprache (DaM), DaZ, bilingual DaM). Im Mittelpunkt der Untersuchung steht die Frage, welche Faktoren die besten Prädiktoren für die Rechtschreiblei-

4 Weitere Informationen unter: <https://exmaralda.org/de/> (13.06.2026)

tung sind und welche quantitativen sowie qualitativen Unterschiede in den Rechtschreibfehlern zwischen den verschiedenen Sprachgruppen bestehen. Die Ergebnisse verdeutlichen, dass die Schulart der wichtigste Prädiktor für die Rechtschreibleistung ist. Nach der Bereinigung der Variablen, die mit der Schulart korrelieren, stellen sich der sozioökonomische Status, der Wortschatz und das Geschlecht als signifikante Einflussfaktoren heraus. Im Gegensatz zu weit verbreiteten Annahmen wirkt sich die Mehrsprachigkeit allein nicht signifikant auf die Rechtschreibleistung aus. In Bezug auf die Fehlerarten finden die Autorinnen heraus, dass Mehrsprachige und Einsprachige größtenteils ähnliche Rechtschreibfehler machen. Dennoch zeigt sich, dass die DaZ-Schüler*innen mehr morphologisch basierte Rechtschreibfehler machen. Außerdem hatten sowohl die DaZ- als auch die bilingualen DaM-Schüler*innen mehr Probleme mit der lautgetreuen Schreibung. Ein signifikanter Unterschied zwischen den Sprachgruppen wird jedoch bei den besonderen Laut-Buchstaben-Beziehungen nicht festgestellt.

Die Längsschnittstudie zur individuellen Rechtschreibentwicklung von Bulut (2018) schlussfolgert ebenfalls, dass Mehrsprachigkeit an sich keinen signifikanten Einfluss auf die Rechtschreibentwicklung von Grundschulkindern hat. Die Untersuchung an rund 900 Erst- und Zweitklässler*innen zeigt zudem, dass interne und externe Faktoren die individuellen Entwicklungswege beeinflussen. Während beispielsweise die phonologische Bewusstheit einen generellen Effekt auf alle Rechtschreibphänomene ausübt, wirken der sozioökonomische Status und kognitive Fähigkeiten spezifischer; so zeigt die kognitive Leistungsfähigkeit den stärksten Einfluss auf die korrekte Schreibung des Schwas (vgl. Bulut, 2018, S. 110). Von besonderer methodischer Bedeutung ist, dass Bulut in dieser Arbeit erstmals die Theorie komplexer dynamischer Systeme (CDST) auf die Rechtschreibentwicklung anwandte. Mit diesem Ansatz kann die individuelle Rechtschreibentwicklung theoretisch modelliert und empirisch gestützt werden. Die Ergebnisse legen nahe, dass die Interaktion sozialer und kognitiver Faktoren nicht zu universellen, vorhersagbaren Entwicklungsstufen führt, sondern vielmehr zur Entstehung individueller Attraktorzustände und Entwicklungsverläufe (vgl. Bulut, 2018, S. 111).

Im Vergleich zu den vielfältigen Untersuchungen über den Orthographieerwerb der mehrsprachigen DaZ-Lernenden bestehen erhebliche Forschungslücken im Bereich des Orthographieerwerbs von chinesischen Deutschlernenden. Es gibt im Kontext Chinesisch als L1 bisher kaum Untersuchungen, die sich ausschließlich mit Orthographiekompetenz oder Orthographieleistung chinesischer Lernender beschäftigen. In den Forschungen, in denen lernersprachliche Abweichungen im schriftlichen Ausdruck chinesischer Deutschlernender qualitativ oder/und quantitativ analysiert werden, werden Rechtschreibfehler nur als eine der Fehlerkategorien betrachtet (Chi, 2016; Ding, 2018; Liu, 2019; Tang, 2003; Z. Wu & Li, 2022), dabei werden die Rechtschreibfehler meist mithilfe von Aufsätzen der Lernenden analysiert. Nur in einem Zeitschriftenbeitrag von Xu, Liu und Yu (2016) werden Diktattexte der Lernenden als Untersuchungsgegenstand betrachtet.

Tang (2003) erstellt ein Untersuchungskorpus aus 33 handschriftlichen Aufsätzen, die von den chinesischen Deutschstudent*innen in der PGG-Prüfung im Jahr 1991 geschrieben wurden. Die linguistische Klassifizierung lernersprachlicher Abweichungen

erfolgt in der vorliegenden Untersuchung anhand einer zweidimensionalen Systematik: Diese orientiert sich einerseits an den betroffenen Normbereichen und andererseits an den Komplexitätsebenen der sprachlichen Ausdrücke. Zu den Normbereichen gehören Schreibung, Grammatik, Semantik, Pragmatik. Zu den Komplexitätsebenen gehören Wort-, Satz- und Textebene. Unter Schreibabweichungen werden die Fehler weiterhin in drei Gruppen geteilt: phonographematische Abweichungen, graphematische Abweichungen und Abweichungen der Zeichensetzung. Insgesamt liegen 146 Schreibabweichungen in seiner Untersuchung vor, davon kommen 88 Abweichungen auf der Wortebene vor, 56 Abweichungen auf der Satzebene, 2 Abweichungen auf der Textebene. Von den gesamten 146 Abweichungen entfallen 43 auf phonographematische, 36 auf graphematische Abweichungen und 67 auf abweichende Zeichensetzung (vgl. Tang, 2003, S. 64). Innerhalb der phonographematischen Schreibabweichungen kommen fehlende Grapheme (55,8 %) am häufigsten vor. Darauf folgen fehlerhafte Wahl von Graphemen (27,9 %), „überflüssige“ Grapheme (11,6 %) und fehlerhafte Graphemfolge (4,7 %). Eine genaue Fehlerliste der phonographematischen Abweichungen wird nach absoluter Fehleranzahl erstellt.

Unter den graphematischen Abweichungen sind Groß- und Kleinschreibungen die häufigsten (52,8 %). Darauf folgen Abweichungen durch falsche Wahl von Allographen (30,6 %) und fehlerhafte Getrennt- und Zusammenschreibung (16,6 %). Abweichungen im Bereich der Zeichensetzung werden nach den betroffenen sprachlichen Ebenen klassifiziert. Dabei kommen Abweichungen im Gebrauch der Satzzeichen am häufigsten (86,3 %) vor. Darauf folgen Abweichungen im Gebrauch von Wortzeichen (13,4 %) und Abweichungen auf Textebene (3 %).

Gleich wie die Untersuchung von Tang für den Prüfungsteil schriftlicher Ausdrücke gilt die PGG-Prüfung auch im Fall von Ding (2018) als Forschungsfeld. In ihrer Arbeit werden die lexikalischen Fehler in 100 Aufsätzen von 2009 bis 2011 analysiert. Diese Fehler ordnet sie vier Kategorien zu: Rechtschreibfehler, falsches Genus des Nomens, idiomatischer Fehler und semantischer Fehler. Unter Rechtschreibfehlern sind die Fehler weiter in vier Untergruppen aufgeteilt, nämlich phonologische Abweichung, graphemische Abweichung, morphologische Abweichung und nicht großgeschriebene Initiale des Wortes. In der Untersuchung von Ding wird jedoch keine statistische Auswertung der Fehler angestellt.

Nach ihrer Analyse der typischen phonologischen Abweichungen in den Lernertexten stellt Ding fest, dass die Fehler wie Weglassung, Verwechslung oder Hinzufügung von Lauten wie *<schlimmesten>, *<Kreditarten>, *<Perzon> meistens auf Achtlosigkeit beim Schreiben zurückzuführen sind. Manche falschen Schreibungen von erfundenen Pluralformen wie **die Englischunterrichts* und die falsche Anwendung der englischen Wörter wie *<efficient>, *<generall> beruhen auf dem Einfluss der ersten Fremdsprache Englisch (vgl. Ding, 2018, 8 ff.). Als graphemische Abweichungen bezeichnet Ding (Ding, 2018) den Fehler, der „die Aussprache des Zielwortes nicht verändert, nur die Form des Wortes“. Hier wird zwar eine Fehlerliste mit deren Kontext aufgelistet, aber eine weitere linguistische Fehlerverteilung findet nicht statt. Bei den aufgelisteten Abweichungen handelt es sich einerseits um die Fremdwortschreibung wie z. B. *<Restauran> (*Restau-*

rant), andererseits um die Kürzen- und Längenmarkierung wie z. B. *<zimlich> (*ziemlich*), *<vertreten> (*vertreten*), *<Staal> (*Stahl*); darüber hinaus werden die Verwechslungen zwischen den Graphemen wie <eu> und <äu> bei *<heufiger> (*häufig*) sowie Fehler bei der Auslautverhärtung *<Weld> (*Welt*) auch in diese Kategorie eingeordnet. Eine ausführliche Interpretation, warum diese Laute häufig von chinesischen Lernenden verwechselt, hinzugefügt oder weggelassen werden, wird in der Untersuchung von Ding nicht durchgeführt.

In der Analyse von Rechtschreibfehlern rechnet Ding in der Kategorie morphologische Abweichungen falsche Deklination, Konjugation und unrichtiges Derivat. In der Untergruppe „nicht großgeschriebene Initiale des Wortes“ werden die Abweichungen in den Lernertexten als Beispiele aufgelistet. Nach Ding kann man solche Fehler aus der Perspektive der negativen Übertragung aus dem Englischen erklären (vgl. Ding, 2018, 33 f.). Die häufig in der Lernaltersprache vorkommende Großschreibung statt Kleinschreibung sowie fehlerhafte Zusammen- und Getrenntschreibung werden aber in den Untersuchungen von Ding nicht in Betracht gezogen.

Anstelle der Anwendung von Prüfungsmaterialien haben viele Forscher*innen selbst Lernertexte gesammelt und ein Lernerkorpus erstellt. Chi (2016) sammelt 90 Aufsätze (Personenbeschreibung und Briefe) von 30 Germanistikstudierenden in einer Klasse an der chinesischen Universität USTB als Lernerkorpus (vgl. Chi, 2016, 9 & 16 f.). Insgesamt werden 597 Fehler im Rahmen ihrer Abschlussarbeit analysiert. Diese Fehler lassen sich drei Kategorien zuordnen: morphologische, syntaktische und textsortenbedingte Fehler. Nach der Auswertung der Fehler kommt Chi zu dem Ergebnis, dass 60,9 % der Fehler dem morphologischen Bereich zuzuordnen sind. 31,6 % der Fehler treten auf der syntaktischen Ebene auf, davon gehören 4,5 % zur Kategorie Textkohärenz und Textkohäsion. Letztlich sind 3 % keiner dieser drei Kategorien zuzuordnen. Bemerkenswert an ihrer Arbeit ist die Unterkategorisierung der Fehler: Chi rechnet semantische Fehler und Rechtschreibfehler gemeinsam zur Kategorie der morphologischen Fehler. Die Rechtschreibfehler teilen sich in weitere Untergruppen auf: falsche Schreibung gleich oder ähnlich klingender Laute; Hinzufügung oder Weglassung von Lauten; kein deutsches Wort wie *<Idea>; falsche Wortbildung wie *<Beispielerweise> und fehlerhafte Groß- und Kleinschreibung. Eine detaillierte statistische Analyse der Fehler in den Untergruppen findet in ihrer Untersuchung nicht statt. Nach der Analyse von solchen Rechtschreibfehlern stellt Chi fest, dass sich die meisten Fehler in dem orthographischen Bereich auf nicht gut beherrschte phonetische Regeln oder auf Nachlässigkeit, Müdigkeit sowie Stress zurückführen lassen. Andere Fehler, „kein deutsches Wort“, lassen sich häufig auf negative Transfereffekte aus der L2 Englisch zurückführen (vgl. Chi, 2016, 20 f.).

M. Liu (2019) untersucht die lexikalischen Fehler in 32 Aufsätzen von fünf chinesischen Schüler*innen, die ihre Schullaufbahn mit dem Abiturziel in Deutschland fortsetzen. Als lexikalische Fehler werden in ihrer Arbeit Rechtschreibfehler und semantische Fehler definiert. Die Rechtschreibfehler werden weiter unterteilt in: Flüchtigkeitsfehler, falsche Schreibung gleich oder ähnlich klingender Laute, Fehler in Groß- und Kleinschreibung, Fehler in Getrennt- und Zusammenschreibung sowie „Sonstiges“.

Von insgesamt 729 Fehlern in 6.012 Wörtern werden 176 Rechtschreibfehler analysiert. Die Versuchspersonen machen am häufigsten Flüchtigkeitsfehler (36,3 % aller orthographischen Fehler), zu denen Auslassungen, Hinzufügungen, Buchstabenvertauschungen und optische Verwechslungen zählten. Darauf folgen Fehler in der Groß- und Kleinschreibung (33,0 %) und die falsche Schreibung von Lauten (21,0 %). Fehler in der Getrennt- und Zusammenschreibung treten mit 8,0 % seltener auf.

Gemäß ihrer Interpretation der orthographischen Fehler führt M. Liu (2019, S. 84) bestimmte Flüchtigkeitsfehler, wie die Weglassung von Endplosiven und die falsche Schreibung aufgrund der Verwechslung von langen und kurzen Vokalen, auf interlinguale Interferenz aus der Muttersprache Chinesisch zurück. Als Grund hierfür nennt sie, dass es im Chinesischen keine Unterscheidung der Vokalquantität gebe. Verwechslungen ähnlich klingender Laute wie [n] und [ŋ] seien hingegen auf Interferenzen des jeweiligen Heimatdialekts zurückzuführen, da diese beiden Laute in südchinesischen Dialekten oft identisch ausgesprochen würden. Fehler in der Groß- und Kleinschreibung sowie in der Getrennt- und Zusammenschreibung begründet Liu primär mit einem Mangel an Kenntnissen der deutschen Orthographieregeln. Darüber hinaus zieht sie auch Interferenzen aus anderen bereits erworbenen Fremdsprachen wie Englisch und Französisch als wichtige Erklärung für diese Fehlerarten in Betracht.

In der Studie von Z. Wu und Li (2022) wird aus dem im Aufbau befindlichen Chinesischen Deutschlerner-Korpus (CDLK) ein Teilkorpus erstellt. Nach seiner Fertigstellung wird das CDLK das erste fehlerannotierte Lernerkorpus für chinesische Deutschlernende mit großen Datenmengen sein, insgesamt werden acht Fehlerkategorien im CDLK annotiert: orthographische Fehler, morphosyntaktische sowie morphologische Fehler, syntaktische Fehler, lexikalische Fehler, semantische Fehler, Logikfehler, Stilfehler und Transferfehler (Y. Li & Wu, 2022b). In der Studie von Z. Wu und Li (2022) werden insgesamt 210 Aufsätze mit 31.701 Wörtern von chinesischen Studierenden des Fachs Germanistik aus dem ersten bis dritten Studienjahr an verschiedenen Universitäten analysiert. Als orthographische Fehler werden in dieser Studie Groß- und Kleinschreibungsfehler, Fehler bei Getrennt- und Zusammenschreibung sowie Wortschreibungsfehler identifiziert. Außerdem werden die Abweichungen bei Präfixen, Suffixen sowie bei anderen Wortbildungselementen (z. B. Fugenelemente) unter den morphologischen Fehlerkategorien als Wortbildungsfehler markiert. Bei der Fehleraufteilung wird offensichtlich, dass morphosyntaktische Fehler am häufigsten im Korpus auftreten, danach folgen semantische Fehler und orthographische Fehler sowie idiomatische Fehler. Unter den orthographischen Fehlern treten die Groß- und Kleinschreibungsfehler am häufigsten im zweiten Studienjahr auf, während sich im dritten Studienjahr ein Rückgang herausstellt. Die Fehlerrate bei den Wortschreibungen nimmt jedoch im Lauf der Zeit zu. Im Allgemeinen treten die Fehler bei der Getrennt- und Zusammenschreibung deutlich seltener auf, und die Fehlerrate bleibt in jeder Lernphase fast unverändert. Z. Wu und Li (vgl. 2022, 133 f.) erklären die vermehrten Fehlererscheinungen der Wortschreibungen mit der Aufmerksamkeitshypothese von W. R. Schmidt (1990): „wenn die Lerner nicht auf die Form achten, kann die Interaktion allein die Sprachkorrektheit nicht verbessern.“ Zudem betonen Wu/Li (vgl. 2022, 143), dass Chinesisch als

Muttersprache einen großen Einfluss auf die lexikalischen Fehler chinesischer Deutschlernender habe. Ein genaues Verhältnis zwischen Muttersprache und orthographischen Fehlererscheinungen wird in der Studie jedoch nicht untersucht.

Im Vergleich zu anderen Untersuchungen der lernersprachlichen Abweichungen werden in der Studie von Xu et al. (2016) 100 Diktatübungen von den Germanistikstudierenden am Ningbo Technologie College in der ersten Studienstufe als Korpusdaten benutzt. Die in den Diktattexten vorkommenden Fehler werden in vier Kategorien eingeteilt, nämlich Wortschreibungsfehler, Verständnisfehler, Grammatikfehler und Rechtschreibfehler. Aus fachwissenschaftlicher Sicht ist diese Kategorisierung jedoch kritisch zu hinterfragen, da sie deskriptive Ebenen mit präsumtiven Fehlerursachen vermengt. So werden beispielsweise Abweichungen wie <Duche> und <praktich> als „Wortschreibungsfehler“ geführt, während phonologische oder morphologische Unsicherheiten wie <vervolgt> (*verfolgt*) aufgrund der vermuteten Ursache unter „Verständnisfehler“ subsumiert werden. Zudem erscheint die Einengung des Begriffs „Rechtschreibfehler“ auf lediglich Groß-Kleinschreibung und Interpunktion als methodisch problematisch, da sie die graphematische Dimension weitgehend ausklammert.

Nach der Analyse wird das Fazit gezogen: „Wortschreibungsfehler“ und „Grammatikfehler“ machen den Hauptteil der gesamten Fehler in den Diktaten aus. Fehler in der Kategorie „Rechtschreibung und Interpunktion“ kommen hingegen relativ wenig vor. Xu et al. (2016) erklären die Wortschreibungsfehler mit unbekanntem Wortschatz, Achtlosigkeit, schwachem Gedächtnis und Verwechslung mit anderen ähnlichen Wörtern. Eine detaillierte Analyse der Fehlerart und Fehlerverteilung findet in der Untersuchung nicht statt.

Durch die Skizzierung des Forschungsstands zur Orthographieleistung und -kompetenz chinesischer Deutschlernender wird eine signifikante Forschungslücke deutlich. Einerseits fehlt es an Studien, die den Erwerb der Orthographie als eigenständiges Forschungsobjekt in den Fokus rücken. Die Mehrheit der bisherigen Untersuchungen behandelt Orthographiefehler als Ganzes, ohne eine detaillierte Differenzierung der Fehlerarten vorzunehmen. Insbesondere eine spezifische Kategorisierung von Orthographiefehlern findet nur selten statt.

In den wenigen Studien, die orthographische Fehler chinesischer DaF-Lernenden nach Fehlerarten klassifizieren, zeigen sich erhebliche Uneinheitlichkeiten sowohl in der Fehlerkategorisierung als auch in den zugrunde liegenden Kriterien. Zudem werden Orthographiefehler selten auf Basis der deutschen orthographischen Prinzipien systematisch kategorisiert und analysiert. Die meisten dieser Studien konzentrieren sich primär auf Fehler der phonographischen Ebene, beispielsweise fehlerhafte Graphem-Phonem-Korrespondenzen, Auslassungen, Hinzufügungen oder falsche Anordnungen von Graphemen.

Ein weiterer Forschungsschwerpunkt sind Fehler auf der syntaktischen Ebene, insbesondere der Groß- und Kleinschreibung sowie der Getrennt- und Zusammenschreibung. Das Phänomen der Vokalquantität und Konsonantenverdopplung – dessen Relevanz im DaZ-Kontext bekannt ist – wird hingegen nur in wenigen Studien explizit für diese Lernergruppe untersucht. Tang (2003) und M. Liu (2019) analysieren

diese Aspekte teilweise in den Fehlerkategorien „graphemische Abweichungen“ oder „falsche Schreibung gleichklingender Laute“. Andere Studien erwähnen solche Fehler lediglich beispielhaft, ohne sie weitergehend zu analysieren. Darüber hinaus finden wesentliche Phänomene der deutschen Orthographie wie die s-Schreibung, die Auslautverhärtung und die Wahrung der morphologischen Konstanz in den genannten Studien kaum oder keine Berücksichtigung. Auch die silbische Schreibung wird nicht erwähnt. Fehler in der Zeichensetzung thematisiert ausschließlich Tang (2003).

Die Erklärungsansätze für die Rechtschreibfehler chinesischer Lernender sind in den referierten Studien heterogen. Einige Autor*innen betonen den sprachlichen Kontrast zwischen der Erstsprache Chinesisch und der Zielsprache Deutsch als primäre Fehlerursache (Ding, 2018; Liu, 2019; Tang, 2003; Z. Wu & Li, 2022). Andere Studien führen den potenziellen Einfluss bereits erworbener Fremdsprachen wie Englisch oder Französisch als Fehlerquelle an (Chi, 2016; Liu, 2019). Andere Untersuchungen sehen wiederum die Hauptursache für die orthographischen Probleme in intralinguistischen Schwierigkeiten des Deutschen selbst (Tang, 2003; Xu et al., 2016), die ein Hauptgrund für die orthographischen Probleme seien. Neben den beschriebenen Unschärfen in der Fehlerklassifikation zeigt sich eine weitere Lücke in der aktuellen Forschungslandschaft: Bisherige Studien konzentrieren sich zumeist auf die deskriptive Analyse der Fehlerphänomene, während externe Lernfaktoren wie Schultyp, Lernzeit oder Lernort nur selten systematisch in die Untersuchung der orthographischen Leistungen chinesischer Lernender einbezogen werden.

Um die spezifischen Charakteristika und die dynamische Entwicklung der orthographischen Kompetenz chinesischer Deutschlernender im Kontext ihres multikompetenten Systems tiefgreifender zu verstehen, ist eine systematische Untersuchung ihrer Rechtschreibfehler unerlässlich. Konkret gilt es dabei zu untersuchen:

1. Mit welchen spezifischen Herausforderungen sind chinesische DaF-Lernende beim Erwerb der deutschen Rechtschreibung konfrontiert, und welche typischen Fehlermuster lassen sich bei ihnen beobachten?
 - Welche konkreten Schwierigkeiten treten in den verschiedenen orthographischen Bereichen (phonographisch, silbisch, morphologisch, syntaktisch und sonstige Fehlerbereiche) auf?
 - Inwieweit sind spezifische Rechtschreibphänomene als Ergebnis der Interaktion zwischen der Erstsprache Chinesisch (L1), den erworbenen Fremdsprachen, insbesondere Englisch, und der Zielsprache Deutsch innerhalb des individuellen, multikompetenten Repertoires der Lernenden zu interpretieren?
2. Welchen Verlauf nimmt die Entwicklung der Rechtschreibkompetenz bei chinesischen DaF-Lernenden in der Anfangsphase des Deutschlerwerbs, und welche Attraktorzustände lassen sich dabei identifizieren?
3. Welche Faktoren beeinflussen die beobachtbare orthographische Leistung chinesischer DaF-Lernenden, und wie ist deren dynamisches Zusammenspiel zu charakterisieren? Welchen Einfluss haben insbesondere der regionale Lernkontext, der Hochschultyp, die Lerndauer und die Art der Aufgabenstellungen auf die Rechtschreibleistung?

3 Methodik

3.1 Lernerkorpora in der Spracherwerbsforschung

In der Spracherwerbsforschung spielt das Lernerkorpus eine bedeutende Rolle, insbesondere wenn es um die Untersuchung von Lernprozessen und Erwerbsphasen geht. Korpora dienen als Grundlage für die systematische Analyse von Sprachdaten und sind unverzichtbar, um Aussagen über bestimmte sprachliche Merkmale zu treffen. Sie bieten die Möglichkeit, empirisch fundierte Erkenntnisse zu gewinnen und die Entwicklung von Lernenden zu verfolgen. In diesem Kapitel wird die korpuslinguistische Methode vorgestellt und ein Überblick über bestehende Lernerkorpora im DaF-Bereich gegeben.

3.1.1 Korpuslinguistische Zugänge

Die korpuslinguistische Methode ist eine empirische Herangehensweise, die auf der Analyse authentischer Sprachdaten basiert, die systematisch gesammelt und in einem Korpus organisiert werden. Ein solches Korpus, wenn es Sprachdaten von Lernenden einer Sprache enthält, wird als Lernerkorpus bezeichnet. Nach Granger, Gilquin und Meunier (2015, S. 1) sind Lernerkorpora „[...] electronic collections of natural or near-natural data produced by foreign or second language (L2) learners and assembled according to explicit design criteria“. Im Rahmen dieser Arbeit wurde ein solches Lernerkorpus erstellt und für die Analyse der orthographischen Kompetenz chinesischer Deutschlernender herangezogen.

Mithilfe der korpuslinguistischen Methode lassen sich typische Lernphänomene und Entwicklungsphasen in der Lernaltersprache systematisch identifizieren, kategorisieren und quantifizieren. Dies führt zu einem tieferen Verständnis der spezifischen Herausforderungen und Lernprozesse der Lernenden. Aus einer CDST-Perspektive bieten Lernerkorpora eine wertvolle Datenbasis, um die für dynamische Systeme charakteristische Nicht-Linearität und Variabilität im L2-Entwicklungsprozess sichtbar zu machen (vgl. Verspoor et al., 2021, S. 172–173). Dabei sind insbesondere longitudinale Lernerkorpora mit einer hohen Dichte an Messpunkten („dense data“) ideal, um individuelle Entwicklungstrajektorien detailliert nachzuzeichnen und die komplexen Interaktionen zwischen verschiedenen sprachlichen Subsystemen über die Zeit zu untersuchen (vgl. Verspoor et al., 2021, 173, 187–188). Die systematische Sammlung von Lernerdaten in einem Korpus bildet somit eine unverzichtbare Grundlage, um den Orthographieverwerb als individuellen Entwicklungsprozess zu untersuchen. Anstatt Fehler lediglich als Abweichungen von einer statischen Norm zu zählen, ermöglicht das Korpus die Analyse erwerbsbedingter Veränderungen im Zeitverlauf. Dabei erlaubt die Verknüpfung

der schriftsprachlichen Korpusdaten mit Metadaten, verschiedene interne und externe Einflussfaktoren systematisch in die Untersuchung einzubeziehen und deren Auswirkung auf die Rechtschreibleistung differenziert zu bewerten.

3.1.2 Bestehende fehlerannotierte DaF-Lernerkorpora

Für die Untersuchung des Erwerbs orthographischer Kompetenz ist die Analyse von Fehlern zentral. Fehlerannotierte Lernerkorpora, die systematisch Abweichungen von der Zielsprache dokumentieren, bieten eine entscheidende empirische Grundlage. Im DaF-Kontext existieren mehrere fehlerannotierte Korpora, die für die vorliegende Arbeit relevante Vergleichsdaten liefern können.

Zu den bekanntesten Lernerkorpora zählt die Falko-Familie⁵ (Hirschmann et al., 2022), eine Sammlung verschiedener Lernerkorpora, die sich durch eine detaillierte Annotation auf mehreren Ebenen auszeichnet. Ein Kernmerkmal ist die Annotation mittels expliziter Zielhypothesen (*Target Hypotheses TH*). Die Formulierung einer solchen Zielhypothese ist für eine systematische und nachvollziehbare Fehleranalyse unerlässlich. Ohne diese explizite Referenz bliebe die Interpretation, was genau als Fehler gewertet wird und wie die korrekte Form lauten sollte, oft implizit und potenziell inkonsistent zwischen verschiedenen Annotierenden (vgl. Lüdeling & Hirschmann, 2015, 136 f.). Im Falko-Projekt wird dieser Ansatz konsequent umgesetzt, wobei die Zielhypothese explizit nicht primär als pädagogische Fehlerkorrektur, sondern als methodischer Normalisierungsschritt dient (vgl. Hirschmann, 2019, 29 f.). Der Zweck ist es, Abweichungen von der zielsprachlichen Norm nach einheitlichen Kriterien strukturell zu erfassen, um sie einer systematischen Analyse zugänglich zu machen. Technisch wird dies in Falko durch eine mehrschichtige Annotation realisiert: Neben der Ebene der Lerneräußerung (word) gibt es mindestens eine Zielhypothesen-Ebene (ZH1), und die Unterschiede dazwischen werden oft auf einer weiteren Ebene (ZH1Diff) durch Editierdistanz-Codes (wie CHA für Änderung, DEL für Löschung, INS für Einfügung) explizit markiert (vgl. Reznicek et al., 2012; Reznicek, Lüdeling & Hirschmann, 2013). Falko enthält Texte fortgeschrittener DaF-Lernender unterschiedlicher Erstsprachen, darunter auch ein Subkorpus mit 20 Texten von Lernenden mit L1 Chinesisch (im Rahmen des Kobalt-DaF-Korpus).

Ein weiteres wichtiges, fehlerannotiertes Korpus ist MERLIN⁶ (MERLIN-Projekt, 2014b), das Lernertexte für Deutsch, Italienisch und Tschechisch zusammenstellte, die zuverlässig den Niveaustufen des Gemeinsamen Europäischen Referenzrahmens (GER) zugeordnet sind. Die Annotation von Abweichungen basiert auch hier auf expliziten Zielhypothesen (TH), wobei sich die Annotation grammatischer und orthographischer Fehler (EA1) auf eine minimale Zielhypothese (TH1) bezieht. Diese stellt eine lediglich orthographisch und grammatisch korrigierte Version dar, bei der nur minimale Eingriffe vorgenommen wurden (vgl. Annotationsschema MERLIN-Projekt, 2014a). Für das Deutsche werden u. a. die orthographischen Phänomene in sieben Kategorien erfasst und

5 Weitere Informationen zum Falko-Projekt unter: <https://www.linguistik.hu-berlin.de/de/institut/professuren/korpuslinguistik/forschung/falko> (22.05.2026).

6 Weitere Informationen zum MERLIN-Projekt unter: <https://www.merlin-platform.eu/#> (22.05.2026).

annotiert: allgemeiner Graphemfehler, Graphemumstellung, Groß-/Kleinschreibung, Wortgrenzen, Abkürzung, Interpunktion und Apostroph. Obwohl das MERLIN-Korpus insgesamt 1.033 Texte von Deutschlernenden enthält, stammen nur 10 Texte von Lernenden mit Chinesisch als L1.

Das Dulko-Korpus⁷ (Beeh et al., 2021) verwendet ein mehrdimensionales System zur expliziten Annotation von Fehlern. Es gibt fünf Hauptfehlerebenen: FehlerOrth (Ebene für orthographische Fehler); FehlerMorph (Ebene für morphologische Fehler); FehlerSyn (Ebene für syntaktische Fehler); FehlerLex (Ebene für lexikalische Fehler); FehlerSem (Ebene für semantische Fehler). Auf der Ebene *FehlerOrth* werden spezifische Kategorien verwendet, um orthographische und Interpunktionsfehler zu klassifizieren. Gemäß dem Dulko-Handbuch (Beeh et al., 2021, S. 33–35) sind dies die folgenden vier Kategorien:

- **GKS:** Kennzeichnet Fehler bei der Groß- und Kleinschreibung.
- **GZS:** Kennzeichnet Fehler bei der Getrennt- und Zusammenschreibung.
- **WS:** Kennzeichnet Fehler bei der Wortschreibung, d. h. wenn in Wörtern Buchstaben fehlen, überflüssig sind oder verwechselt werden.
- **ZS:** Kennzeichnet Fehler bei der Zeichensetzung (Interpunktion), sowohl auf Satz- als auch auf Wortebene (z. B. bei Abkürzungen).

Es ist anzumerken, dass die Architektur von Dulko durch die Verwendung multipler Zielhypothesenebenen durchaus in der Lage ist, mehrere Abweichungen innerhalb eines Wortes isoliert abzubilden. Dennoch operiert die Annotation auf einer kategorialen Makroebene. Selbst wenn verschiedene Abweichungen auf getrennten Ebenen erfasst werden, erfolgt die Zuweisung oft über pauschale Tags wie *WS*. So wird beispielsweise die Fehlschreibung <Bankkonnto> (DulkoEssay-v.1.0-Entlohnung 27) zwar mit dem Tag *WS* annotiert, eine weiterführende sublexikalische Differenzierung unterbleibt jedoch. Eine detaillierte Klassifizierung dieses Fehlers geht somit nicht unmittelbar aus dem Tag hervor, wie etwa im Fall der Schreibung <Bankkonnto>, die linguistisch als Verstoß gegen silbische Prinzipien zu werten wäre.

Im Rahmen des Dulko-Projekts wurde eine Reihe von spezialisierten Software-Werkzeugen von Andreas Nolda entwickelt: die sogenannten „Dulko-Tools“ für den EXMARaLDA-Partitur-Editor (vgl. Nolda, 2024). Diese Toolsammlung ermöglicht eine effiziente und standardisierte semiautomatische Annotation von Lernertexten, inklusive Tokenisierung, POS-Tagging, Lemmatisierung sowie der Erstellung und Annotation von Zielhypothesen und Fehlern auf mehreren Ebenen (Nolda, 2024, S. 228). Aufgrund ihrer Praktikabilität und des Open-Source-Ansatzes wurden diese Dulko-Tools auch von anderen Lernerkorpusprojekten wie NaLeKo (Hirschmann, Binanzer & Langlotz, 2023) aufgegriffen und nachgenutzt (vgl. Nolda, 2024, S. 225). In Anerkennung ihrer Bedeutung sind diese Tools seit 2024 offiziell in der neuen Version von EXMARaLDA-Partitur-Editor (1.8.1) integriert und werden dort weiterentwickelt (vgl. Nolda, 2024, 225 f.).

7 Weitere Informationen zum Dulko-Projekt unter: <https://www.ids-mannheim.de/gra/projekte/deutung/dulko/> (13.06.2026).

Für den Aufbau und die Aufbereitung des DeChiLko-Korpus in der vorliegenden Arbeit wurde ebenfalls auf diese etablierten Werkzeuge zurückgegriffen, insbesondere auf die Annotationsfunktionen in EXMARaLDA (Dulko)⁸ sowie auf das von Nolda entwickelte Korpusbausystem *makeDulko*⁹ für die Konvertierung der Daten. Da die spezifischen Anforderungen dieser Untersuchung jedoch teilweise von den Standardprozeduren abwichen, wurden bestimmte Funktionalitäten und Prozesse an die eigenen Forschungsbedürfnisse angepasst. Eine detaillierte Beschreibung der im Rahmen dieser Arbeit vorgenommenen Prozesse und des konkreten methodischen Vorgehens bei der Korpusaufbereitung und Annotation erfolgt in den nachfolgenden Abschnitten.

Obwohl die genannten Korpora wie Falko, Kobalt-DaF, MERLIN bereits Lerner-texte chinesischer Deutschlerner anbieten, ist ihre Anzahl sehr beschränkt. Die Entwicklung der Lernerkorpora für chinesische Deutschlerner reicht bis nach 2008. In diesem Jahr wurde an der Fremdsprachenuniversität Beijing (BFSU) unter der Leitung von Prof. Qian Minru das „Korpus der Inter Sprache (Deutsch) der chinesischen Germanistikstudenten“ (KICG) zusammengestellt. Dieses KICG, das Aufsätze aus Semesterprüfungen sowie den landesweiten PGG- (Prüfung für Grundstudium der Germanistik) und PGH-Prüfungen (Prüfung für Hochstudium der Germanistik) umfasste, gilt als das erste deutschsprachige Lernerkorpus in der chinesischen Germanistik und diente als Grundlage für viele wissenschaftliche Abschlussarbeiten (vgl. Y. Li & Ge, 2019, S. 193). Jedoch wurden Lernertexte im KICG nicht annotiert und sind nicht öffentlich zugänglich.

Angesichts dieser Lücke wurde das Chinesische Deutschlerner-Korpus (CDLK) unter der Leitung von Prof. Dr. Yuan Li an der Zhejiang-Universität initiiert (Y. Li & Wu, 2022a). CDLK zielt darauf ab, den Schriftspracherwerb chinesischer Lernender von der Schule bis zur Universität umfassend zu dokumentieren und analysierbar zu machen. Es ist zu beachten, dass sich das CDLK-Projekt noch in einer aktiven Aufbau- und Entwicklungsphase befindet; ein umfassendes Handbuch ist zum aktuellen Zeitpunkt (Sep. 2025) noch nicht veröffentlicht, und ein öffentlicher Zugang zum Korpus besteht ebenfalls noch nicht.

Ein zentraler Aspekt des CDLK ist die detaillierte Fehlerannotation. Die hierfür entwickelten Fehlerkategorien bauen auf der Taxonomie des Dulko-Projekts auf (bereitgestellt durch H. Hirschmann), wurden jedoch spezifisch für die Äußerungen chinesischer Deutschlerner „kultursensibel und -spezifisch“ angepasst, differenziert und erweitert (vgl. Y. Li & Wu, 2022a, S. 227). Die Annotation erfasst Abweichungen auf Ebenen wie Orthographie, Morphosyntax, Morphologie, Syntax, Lexik, Semantik, Logik, Stil und Transfer (vgl. Y. Li & Wu, 2022a). Die Fehlschreibungen werden im orthographischen Bereich in drei Kategorien eingeordnet: Groß-/Kleinschreibung, Getrennt-/Zusammenschreibung und Wortschreibung (vgl. Z. Wu & Li, 2022, S. 123). Ähnlich wie im

8 Zum Zeitpunkt der Erstellung und Annotation des DeChiLko-Korpus für die vorliegende Arbeit (vor Dezember 2024) waren die Dulko-Tools noch nicht offiziell in den EXMARaLDA-Partitur-Editor integriert. Daher wurde für die Bearbeitung und Annotation der Lernertexte die Version EXMARaLDA (Dulko) verwendet. Weitere Informationen unter: <https://sr.ht/~nolda/exmaralda-dulko/> (13.06.2026)

9 Das Build-System *Makedulko*, entwickelt von Andreas Nolda, ermöglicht die Konvertierung von Lernerkorpora in ein für ANNIS nutzbares Format. Weitere Informationen unter: <https://sr.ht/~nolda/makedulko/> (13.06.2026)

Dulko-Korpus wird, wenn innerhalb eines Wortes mehrere Abweichungen auftreten, die derselben orthographischen Fehlerkategorie (z. B. der allgemeinen Wortschreibung) zuzuordnen wären, das Wort in der Regel nur einmalig unter dieser Kategorie erfasst. Eine weitergehende sublexikalische Klassifikation von Wortschreibungsfehlern ist somit im CDLK nur eingeschränkt möglich.

Zusammenfassend lässt sich festhalten, dass es notwendig ist, ein spezifisches, auf die orthographische Kompetenz chinesischer Deutschlernender zugeschnittenes Korpus zu erstellen und dabei sowohl auf etablierte Verfahren als auch auf eigene Anpassungen zurückzugreifen. Im folgenden Kapitel wird nun detailliert auf die Erstellung dieses eigenen Lernerkorpus DeChiLKo eingegangen.

3.2 Erstellung des Lernerkorpus DeChiLKo

Wie bereits in Kapitel 3.1.2 dargestellt wurde, stand bis zur Entstehung der vorliegenden Arbeit noch kein frei zugängliches Lernerkorpus spezifisch für chinesische Deutschlernende mit großer Datenmenge zur Verfügung. Um die orthographische Kompetenz der chinesischen Deutschlernenden zu erforschen, stehen die eigenständige Erstellung des Lernerkorpus DeChiLKo und die Annotationen der orthographischen Abweichungen in den Lernertexten bei der Methode dieser Arbeit im Zentrum. Abbildung 5 bietet einen Überblick über die einzelnen Arbeitsschritte bei der Korpuserstellung. Neben der Transkription der als Bilddateien vorliegenden handschriftlichen Texte und dem Import digitaler Textdateien wird in Abbildung 5 auch die Abfolge der Annotationen veranschaulicht.

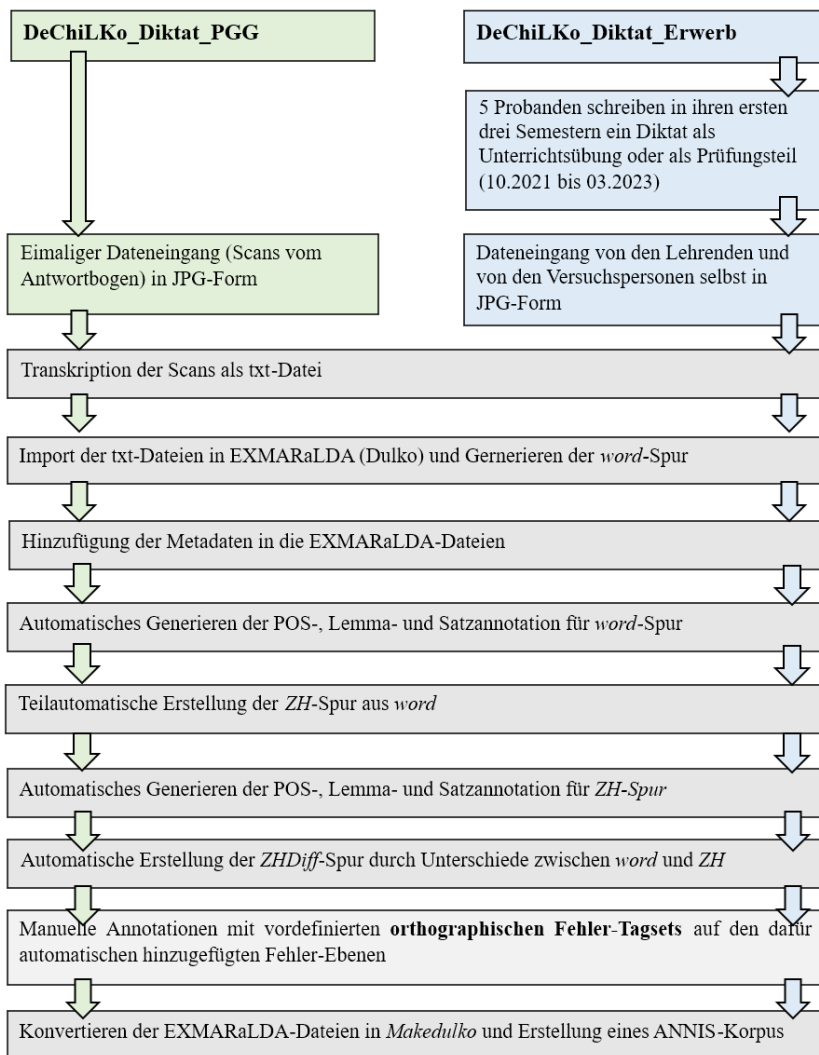


Abbildung 5: Überblick über die einzelnen Arbeitsschritte der Korpuserstellung

3.2.1 Datenerhebung

3.2.1.1 Korpusdaten: Von „Aufsatz“ zu „Diktat“

Bei der Erstellung eines Korpus ist die sorgfältige Auswahl der Datenquelle ein entscheidender erster Schritt (vgl. Hirschmann, 2019, 19 f.). Die Daten in existierenden schriftlichen Lernerkorpora bestehen hauptsächlich aus Textproduktionen der Lernenden; darunter fallen verschiedene Textsorten wie Tagebücher, wissenschaftliche Arbeiten, Aufsätze, schriftlicher Ausdruck zu einem bestimmten Thema, Bildbeschreibungen usw.

In Anlehnung an bekannte Lernerkorpora wie das Falko-Korpus oder das Kobalt-DaF-Korpus, die Essays sowie argumentative Aufsätze der DaF-Lernenden als Korpusdaten beinhalten (vgl. Reznicek et al., 2012), wurde anfangs für das Korpusdesign des DeChiLKo die Textsorte „Aufsatz“ in Betracht gezogen. In der Probeanalyse wurden hundert Aufsätze zum Thema „Übergewicht und Adipositas“ untersucht. Diese Aufsätze gehörten zum Teil „Schriftlicher Ausdruck“ der Prüfung Germanistik Grundstudium (PGG) im Jahr 2019. Die Lernenden sollten unter den Prüfungsbedingungen einen zusammenhängenden Text mit einem Umfang von 150 Wörtern verfassen. Da die sprachlichen Inputs von Fremdsprachenlernenden ohne direktes zielsprachliches Umfeld meist aus Lehrwerken sowie von den Lehrkräften im Unterricht stammen, fiel bei der Analyse auf, dass kein breites Spektrum an Wörtern in diesen Texten verwendet wurde, was möglicherweise auf die thematische Eingrenzung und die Prüfungssituation zurückzuführen ist. Des Weiteren waren in dieser begrenzten und kontrollierten Wortwahl nur wenige orthographische Abweichungen zu finden. Die auftretenden Rechtschreibfehler konzentrierten sich meist auf Wörter, die einerseits den Lernenden weniger geläufig waren, andererseits für die thematische Bearbeitung unvermeidbar erschienen, wie z. B. der Fachbegriff „Adipositas“. Eine Untersuchung der orthographischen Kompetenz chinesischer Deutschlernender anhand solcher Aufsatzkorpora, in denen potenziell nur wenige, oft gleichartige orthographische Abweichungen vorkommen, erschien daher als nicht ausreichend aussagekräftig. Aus diesem Grund wurde für das DeChiLKo auf die Textsorte „Aufsatz“ als primäre Datenquelle für die Orthographieanalyse verzichtet, und stattdessen wurde auf eine im Fremdsprachenunterricht häufig verwendete Übungsform zurückgegriffen: das Diktat.

Diese Aufgabenform ist in der Fremdsprachendidaktik nicht unumstritten. Kritisiert wird unter anderem, dass emotionale und andere lernerexterne Faktoren die Sprachproduktion im Diktat stark beeinflussen können und die Testsituation eine hohe Belastung für die Schreibenden darstellen kann, was zu einer erhöhten Anzahl an Performanzfehlern führen kann (vgl. Birkel, 2009, 5, 10). Dieser potenziellen Problematik war sich die Forscherin durchaus bewusst. Dennoch sprachen mehrere Gründe dafür, Diktattexte als empirische Basis zu nutzen. Zum einen war die einfache Verfügbarkeit solcher Daten ein pragmatischer Faktor. Zum anderen stellt das Diktat im Vergleich zu freien Schreibaufgaben höhere Anforderungen an die kognitiven Ressourcen der Schreibenden, da es ein simultanes Zusammenspiel mehrerer sprachlicher Fertigkeiten erfordert: Hören, Verstehen, Speichern, Schreiben sowie internalisiertes Regelwissen der Morphosyntax und Graphem-Phonem-Korrespondenzen. In diesem Sinne eignen sich Diktate trotz der genannten Kritikpunkte als aufschlussreiche Datenquelle für orthographische Fragestellungen (vgl. Hans-Bianchi, 2007, S. 85; Hans-Bianchi & Katelhön, 2010, S. 36).

Deshalb wurde trotz aller Kritik am Diktat die Entscheidung getroffen, Diktattexte von den chinesischen Deutschlernenden als Korpusdaten zu verwenden.

3.2.1.2 Datenquelle für das Prüfungskorpus: PGG-Prüfung

Alle Diktate im Prüfungskorpus sind unter PGG-Prüfungsbedingungen entstanden. Die PGG-Prüfung ist eine jährlich im Juni stattfindende landesweite Sprachprüfung, welche Studierende in Studiengängen Deutsch/Germanistik in China zum Ende des zweiten Studienjahrs ablegen. Es handelt sich dabei um eine schriftliche Prüfung zur Feststellung der sprachlichen Leistung, die das Bildungsministerium der Volksrepublik China einheitlich veranstaltet. Die Prüfung findet zeitgleich an verschiedenen Hochschulorten statt und steht unter der Aufsicht der örtlichen Prüfungskommission, die sich aus Vertretern der betreffenden Hochschulen zusammensetzt (vgl. Tang, 2003, S. 41). Die PGG-Prüfung umfasst ungefähr das Sprachniveau B1 nach dem europäischen Referenzrahmen (vgl. Zhao, 2020: 55). Die Prüfung dauert insgesamt 170 Min. und besteht aus fünf Teilen, nämlich Diktat (15 Min.), Hörverstehen (20 Min.), Leseverstehen (50 Min.), Grammatik und Wortschatz (50 Min.) sowie Schriftlicher Ausdruck (15 Min.); in der folgenden Tabelle sind die detaillierten Informationen der jeweiligen Prüfungsteile dargestellt.

Tabelle 11: Prüfungsordnungen für das Germanistik-Grundstudium (Quelle: Anleitungskomitee für den Fremdsprachenunterricht an Hochschulen des chinesischen Bildungsministeriums, 2013, S. 14)

Prüfungsteile	Aufgaben	Bezeichnung der Prüfungsteile	Aufgabentypen	Anzahl der Aufgaben	Anzahl der Punkte	Gesamte Punktzahl	Vorgesehene Prüfzeit (Minuten)
I	1–8	Diktat	Lücken füllen	8	4	10	15
			Textdiktat	1	6		
II	1–6	Hörverstehen	Mini-Dialoge	6	6	20	20
	7–20		Texte mittlerer Länge	14	14		
III	21–24	Leseverstehen	Matching Options	5	5	20	50
	25–40		Kurze und mittellange Texte	15	15		
IV	41–50	Grammatik und Wortschatz	Single-Choice	10	5	35	45
	51–60		Cloze-Test	10	5		
	61–70		Cloze-Test	10	5		
	71–90		Lücken füllen	20	10		
	91–95		Satzumschreibung	5	5		
	96–100		Satzkorrektur	5	5		

(Fortsetzung Tabelle 11)

Prüfungs- teile	Auf- gaben	Bezeichnung der Prüfungs- teile	Aufgabentypen	Anzahl der Auf- gaben	Anzahl der Punkte	Gesamte Punktzahl	Vorgesehene Prüfzeit (Minuten)
V		Schriftlicher Ausdruck	Text mit The- menvorgabe verfassen	1	15	15	40
Insgesamt					100	100	170

Die Durchführung des Diktats innerhalb der PGG-Prüfung erfolgt über eine landesweit einheitliche Audioaufnahme. Die Audioaufnahme wurde von einem chinesischen Germanistikdozenten gesprochen. Dies gewährleistet einen kontrollierten, am zielsprachigen Standard orientierten auditiven Stimulus, der für alle Prüfungsstandorte ein identisches Sprechtempo sowie eine einheitliche Artikulation garantiert.

Der Text wird dabei in drei Phasen präsentiert: Zunächst erfolgt ein einmaliges Vorlesen des gesamten Textes in normaler Sprechgeschwindigkeit, gefolgt von einer satzweisen Wiedergabe mit ausreichenden Pausen für die Niederschrift. Abschließend wird der Text zur Selbstkontrolle nochmals im Ganzen vorgelesen. Ein individuelles Nachfragen oder eine Beeinflussung des Tempos durch die Teilnehmenden ist im Rahmen dieses Prüfungssettings nicht vorgesehen.

Die PGG ist eine staatlich organisierte Prüfung, die ein breites Spektrum an Daten bietet, weil sie nicht nur einzelne Hochschulen repräsentiert, sondern die Gesamtheit der Studierenden im Fach Germanistik an chinesischen Hochschulen umfassen kann. Die teilnehmenden Institutionen variieren von Universitäten, Fremdsprachenuniversitäten und technischen Hochschulen bis zu Pädagogischen Hochschulen. Um eine möglichst breite Abdeckung zu gewährleisten, erfolgte die Auswahl der 20 Hochschulen nach dem Verfahren der geschichteten gezielten Auswahl (*stratified purposive sampling*) auf Basis zweier Kriterien: der offiziellen Hochschulkategorie (A–E) sowie der geographischen Region. Hinsichtlich der Hochschulkategorie sind alle fünf Kategorien vertreten. Hinsichtlich der geographischen Verteilung deckt die Stichprobe alle sechs Großregionen Chinas ab: Nordchina, Ostchina, Zentral- und Südchina, Südwestchina, Nordostchina sowie Nordwestchina (siehe Tabelle 12; detaillierte Informationen der ausgewählten Hochschulen für das DeChiLKo-Prüfungskorpus siehe Anhang 9.1). Die Auswahl erhebt keinen Anspruch auf statistische Repräsentativität im strengen Sinne, da die Anteile der einzelnen Kategorien und Regionen nicht proportional zur Gesamtheit der über 100 an der PGG beteiligten Institutionen sind. Sie zielt jedoch darauf ab, die institutionelle und regionale Heterogenität des chinesischen Germanistikstudiums systematisch abzubilden.

Tabelle 12: Kategorisierung der chinesischen Hochschulen (Quelle: In Anlehnung an Deutscher Akademischer Austauschdienst, 2019; Goldberger, 2017; Marioulas & Wu, 2015)

Kategorien	Hochschulen	Eigenschaften
A	Öffentliche Universitäten ersten Rangs	• unterstehen in der Mehrzahl direkt dem Bildungsministerium oder anderen zentralen Regierungsstellen • Hochschulen in den Förderprogrammen „985“ und „211“ • nehmen die Studenten mit den höchsten Punktzahlen in der Hochschuleaufnahmeprüfung (Gaokao) auf
B	Öffentliche Universitäten zweiten Rangs	• Provinzuniversitäten • nehmen in der zweiten Zulassungsrunde die Schulabsolventen auf
C	Universitäten zweiten Rangs	• private Universitäten • verleihen Bachelor-Abschlüsse
D	Universitäten dritten Rangs und unabhängige Colleges (An-Institute)	• private Ausgründungen staatlicher oder privater Hochschulen, die den Namen der renommierten Universitäten tragen • niedrigere Zulassungsvoraussetzungen, aber höhere Studiengebühren
E	Hochschulen mit dreijährigen, anwendungsorientierten Studiengängen	• verleihen keinen akademischen Abschluss, sondern berufsqualifizierende Abschlüsse

Nach der Auswahl der teilnehmenden Hochschulen wurden pro Institution jeweils zehn Lernertexte im Wege einer Zufallsstichprobe für das Korpus ausgewählt. Diese Diktattexte wurden im Anschluss für die weitere Korpuserstellung aufbereitet.

3.2.1.3 Datenerhebung für das Erwerbskorpus

Die Daten des Erwerbskorpus stammen von fünf chinesischen Germanistikstudierenden der Anhui-Universität und wurden vom Wintersemester 2021/22 bis zum Wintersemester 2022/23 erhoben, d. h. während der ersten drei Semester ihres Bachelorstudiums. Es handelt sich hierbei um eine longitudinale Datenerhebung. Das Korpus umfasst insgesamt 140 Texte: 15 Diktattexte aus Semesterprüfungen und 125 Diktattexte, die als Unterrichtsübungen handschriftlich erstellt wurden.

Im Gegensatz zum Prüfungskorpus war die Datenquelle im Erwerbskorpus unterschiedlich gestaltet: Der auditive Input erfolgte entweder durch das Vorlesen der jeweiligen chinesischen DaF-Lehrkraft vor Ort oder durch den Einsatz von standardisierten MP3-Audiodateien (beispielsweise aus Lehrwerken oder Übungsbüchern). Unabhängig von der Quelle wurde der Diktattext wie in der PGG-Prüfung dreimal vorgelesen. Während des Unterrichts und der Semesterprüfung war der Einsatz von Hilfsmitteln wie Wörterbüchern oder Smartphones nicht gestattet.

Angaben über die Metadaten in Bezug auf die fünf chinesischen Germanistikstudierenden und die von ihnen verfassten Lernertexte sind zunächst in der folgenden Tabelle dargestellt.

Tabelle 13: Metadaten der Lernenden im Erwerbskorpus

Name	Geschlecht	Anzahl der Lernertexte	Summe der Tokens
LZX	männlich	25	2.680
LWY	weiblich	25	2.575
YYT	weiblich	28	2.663
ZXY	weiblich	31	2.991
ZXX	weiblich	31	3.167

3.2.2 Datenaufbereitung

3.2.2.1 Transkription der Lernertexte ([word]-Ebene)

Da alle Lernertexte handschriftlich verfasst wurden, bestand der erste Schritt der Datenaufbereitung in der Transkription der Diktattexte. Diese wurden in TXT-Dateien übertragen, wobei das zentrale Konzept der Transkription voraussetzt, die Handschrift der Lernenden möglichst originalgetreu in digitaler Form abzubilden. Alle Zeichen, die mit einer Tastatur darstellbar sind, wurden transkribiert. Ausgenommen sind Symbole, die Korrekturen der Lernenden markieren, wie z. B. Löschungen, Durchstreichungen, Absatzmarken oder Vertauschungspfeile. Gelöschte oder durchgestrichene Wörter oder Sätze wurden nicht transkribiert, ebenso wenig wie Vertauschungs- und Einfügemarierungen. Die als vertauscht markierten Abfolgen wurden in der gewünschten Reihenfolge wiedergegeben.

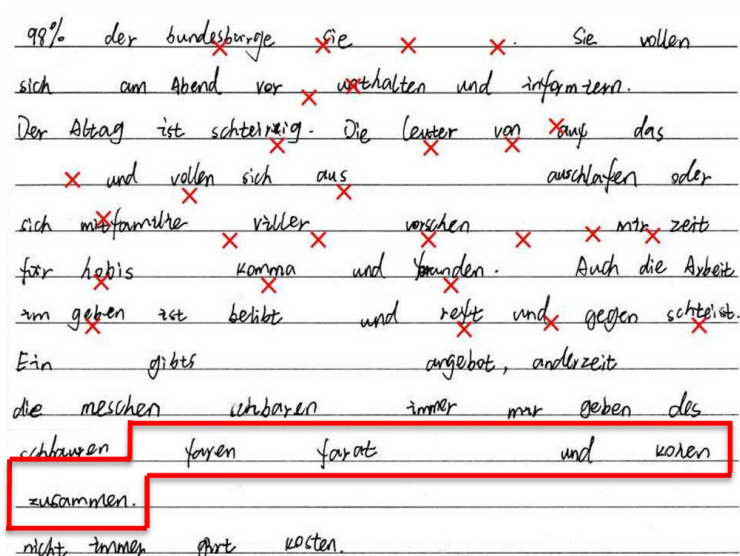


Abbildung 6: Exemplarischer Textausschnitt im Subkorpus DeChiLKo-Prüfung 2019 (Quelle: DeChiLKo_C_NW_2019_D_007)

Die transkribierten Lernertexte wurden für die nachfolgende Annotation in EXMARaLDA (Dulko) (Nolda, 2019) in EXMARaLDA (Dulko) als [Word]-Spur importiert. Abbildung 6 illustriert einen exemplarischen Textausschnitt aus dem Subkorpus DeChiLko-Prüfung 2019. Der rot umrandete Textteil wird für die nachfolgenden Erläuterungen zur Annotationsphase in DeChiLko verwendet.

In der Diktataufgabe entstanden alle Lernertexte auf der Grundlage des vorgelesenen Kurztexts, in diesem Sinne spielen das Sprechtempo, die Betonung sowie die Artikulationsweise der Vorleserin oder des Vorlesers eine wichtige Rolle bei der Produktion des Diktattexts (vgl. Hans-Bianchi, 2007, S. 88). Um die phonetisch-phonologischen Aspekte des auditiven Stimulus präzise darzustellen, wurden die Diktat-Audiodateien in zwei Verfahren transkribiert: nach dem Gesprächsanalytischen Transkriptionssystem 2 (GAT2) (Selting et al., 2009) für die prosodischen Merkmale sowie nach den Konventionen des Internationalen Phonetischen Alphabets (IPA) für die lautliche Realisierung.

Nach GAT2 werden Ziffern und Zahlwörter ausgeschrieben, wie z. B. <neunundachtzig> (vgl. Selting et al., 2009, S. 363). Außerdem wird z. B. der Fokusakzent bei <beLOhnungen> durch Großbuchstaben notiert (vgl. Selting et al., 2009, S. 371). In der IPA-Transkription liegt der Fokus auf phonetisch-phonologischen Merkmalen, Phänomene wie Aspiration bei /t^h/, Haupt- und Nebenbetonung sowie Tonhöhenbewegung lassen sich an dieser Stelle besonders markieren.

Tabelle 14: Ausschnitt der Transkription für Diktataudio

Trans_GAT2	durch	verschiedene	ARten	von	beLOhnungen (.)	WERden
Trans_IPA	ˈdʊəç	fʁɛˈʃiːdənə	ˈʔaːRtn	ˈfɔn	bəˈloːnʊŋən (.)	ˈveːRdən
Word	durch	verschiedene	Arten	von	Belohnungen	werden
Trans_GAT2	die	kinder	zum	WElterlernen	motiviert	punkt (.)
Trans_IPA	diː	ˈkɪndɐ	ˈt͡sʊm	ˈvaɪtɐˈlɛəənən	motiˈviːɐ̯tʰ	ˈpʊŋktʰ(.)
Word	die	Kinder	zum	Weiterlernen	motiviert	.

Die Diktat-Audiodateien für die Lernertexte im Prüfungskorpus wurden in EXMARaLDA transkribiert, und die Transkripte sind als Anhang beigefügt worden (siehe Anhang 8.2). Für das Erwerbskorpus war dies jedoch nicht durchgängig möglich. Aufgrund der spontanen und flexiblen Gestaltung der Unterrichtssituationen lagen nicht für alle Diktateinheiten Audioaufzeichnungen vor. Eine verlässliche Transkription war deshalb nicht realisierbar.

3.2.2.2 Metadaten

Zu allen Texten im DeChiLko-Korpus wurden Metadaten zum Korpusdesign, den Texten und zu den Autor/-innen erhoben. Die Metadaten orientieren sich an dem von Granger und Paquot (2017) vorgeschlagenen Standard.

Nach Import der transkribierten Korpusdaten in EXMARaLDA (Dulko) (Nolda, 2019) wurden mit Hilfe des Transformationsszenarios „Meta“ die Metavariablen au-

tomatisch aus dem Dulko-Template *dulko.template.exb* eingefügt. Alle Metadaten sind auch in den verfügbaren ANNIS-Datenbanken eingepflegt und können dort bei der Korpusuche (siehe Kapitel 3.3) ausgeschöpft werden.

Die Metadaten gliedern sich wie folgt:

- Korpus (Korpusdesign, Informationen zur Annotation)
- Text (Information zur Erhebung, Bewertung und Transkription des Textes)
- Autor/-in (personenbezogene Informationen)

Für das Erwerbskorpus wurden die persönlichen Informationen zu den fünf Versuchspersonen sowie Angaben zu ihrer Sprachlernbiografie mithilfe eines Fragebogens gesammelt. Die Erhebung detaillierter Metadaten über die Teilnehmenden der landesweiten PGG-Prüfung erwies sich hingegen als schwierig. Daher fokussieren sich die Metadaten des DeChiLKO-Prüfungskorpus vorrangig auf text- und korpusbezogene Angaben, während die Informationen über die Lernenden lediglich begrenzt verfügbar sind (eine vollständige Auflistung der verfügbaren Metadaten siehe Anhang 8.3).

3.2.3 Annotation

Die Annotationsphase wurde grundsätzlich in zwei Schritte unterteilt. Im ersten Schritt erfolgte die automatische Annotation der Lemmata, Wortarten¹⁰ und Satzspannen für den Lernertext und die Zielhypothese mithilfe des Transformationsszenarios in EXMARaLDA (Dulko).

[word]	faren	farat	und	kohen	zusammen	.	
[S]							
[pos]	NN	NN	KON	ADJD	PTKVZ	\$.	
[lemma]	faren	farat	und	kohen	zusammen	.	
[ZH]	fahren	Fahrrad	und	kochen	zusammen	.	F
[ZHDiff]	CHA	CHA		CHA			ll
[ZHS]							s
[ZHpos]	VVFIN	NN	KON	VVFIN	PTKVZ	\$.	h
[ZHlemma]	fahren	Fahrrad	und	kochen	zusammen	.	F
[Graph_Fehler]	f	a	ren	f	a	ra	t
[Graph_ZH]	f	ah	ren	F	ah	r	ra

Abbildung 7: Screenshot der ersten Annotationen des Lernertexts (C_NW_2019_D_007) mit Satzspannen (Zeile [S]), Wortarten (Zeile [pos]) und Lemmata ([lemma]) sowie der zusätzlichen Spuren [ZH], [ZHDiff], erneutem Tagging von Satzspannen, Wortarten und Lemmata ([ZHS], [ZHpos] und [ZHlemma]) in EXMARaLDA (Dulko)

Als Zielhypothese ([ZH]) dienen dabei die Musterlösungen der jeweiligen Diktate. Ein weiteres Transformationsszenario „ZHDiff-Spur“ in EXMARaLDA (Dulko) erkennt automatisch Abweichungen zwischen Tokens der Zielhypothese ([ZH]) und des Lernertexts ([word]). Diese Differenzen werden mit den vorgesehenen Abweichungstags auf der [ZHDiff]-Ebene markiert: CHA (verändertes Token), INS (hinzugefügtes Token), DEL (überflüssiges Token), SPLIT (geteiltes Token), MERGE (zusammengesetztes

10 Annotieren von Wortarten ist eine integrierte Funktion im Transformationsszenario in EXMARaLDA (Dulko) mithilfe des *TreeTagger* nach Schmid (1994).

Token), MOVS und MOVT (Token mit anderer Position) (vgl. Beeh et al., 2021; Nolda, 2019; Reznicek et al., 2012). Der in Abbildung 7 geschilderte Aufbau illustriert die ersten Annotationen eines Lernertexts in EXMARaLDA (Dulko).

Die ZHDiff-Ebene stellt eine der wichtigsten Annotationsebenen im DeChiLko dar, weil sie Abweichungen zwischen den von den Lernenden aufgeschriebenen Diktattexten und dem Lösungstext (Zielhypothese) auf Tokenebene erfasst. Eine Analyse der Annotationsverteilung auf dieser Ebene bietet einen ersten Überblick über die orthographische Leistung der Lernenden: Wenn die Annotation auf der [ZHDiff]-Ebene „CHA (change)“ lautet, handelt es sich mit hoher Wahrscheinlichkeit um einen Schreibfehler (inkl. Wortschreibung, Zeichensetzung, Groß- und Kleinschreibung usw.). Die Annotationen „SPLIT“ oder „MERGE“ hingegen repräsentieren typischerweise Fehler im Bereich der Getrennt- oder Zusammenschreibung.

Für eine detaillierte Bewertung der Entwicklung der Rechtschreibkompetenz ist jedoch ein spezifischeres Annotationsschema erforderlich, das die besonderen Eigenschaften der falsch geschriebenen Wörter differenziert erfasst.

Es existieren bereits zahlreiche Annotationsschemata für das Deutsche, die vor allem zur Bewertung der Rechtschreibkompetenz von Kindern mit Deutsch als Erstsprache entwickelt wurden (siehe 2.1.3). Diese Schemata basieren häufig auf Modellen des orthographischen Erwerbs und ordnen Fehler den jeweiligen Erwerbsphasen zu, statt systematisch graphematisch fundierte Kategorien zu verwenden (siehe 2.1.2). Diagnostische Ansätze wie die Hamburger Schreib-Probe (HSP) (May, 2013) und die Oldenburger Fehleranalyse (OLFA) (Thomé & Thomé, 2020) betrachten Fehler zwar im Kontext von Erwerbsphasen, erlauben jedoch selten unmittelbare Rückschlüsse auf die Systematik des deutschen Schriftsystems.

Die Aachener Förderdiagnostische Rechtschreibanalyse (AFRA) (Herné & Naumann, 2016) basiert weitgehend auf graphematischen Prinzipien, weist jedoch teilweise eine fehlende Transparenz in der Zuordnung von Fehlern auf. Beispielsweise wird die Fehlschreibung *<Warheit> (*Wahrheit*) als Fehler bei der Morphem-Differenzierung eingeordnet (Herné & Naumann, 2016, S. 29), obwohl sie ebenso auf eine fehlerhafte Vokalquantitätsmarkierung zurückzuführen sein könnte.

Ein stärker systemorientiertes Annotationsschema findet sich bei Thelen (2010), das das graphematische System in hohem Maße berücksichtigt. Das Schema nimmt die Silbe als zentrale Einheit und unterscheidet systematisch zwischen phonologischen und morphologischen Schreibungen. Es kodiert, ob Silbenanfangsrand, -kern oder -endrand sowie spezifische orthographische Phänomene (z. B. Konsonantendoppelung, markierte Vokalquantität) korrekt geschrieben wurden.

Das Annotationsschema von DeChiLko setzt an den Ansätzen von Thomé (1987), Thelen (2010), Herné und Naumann (2016) sowie Thomé und Thomé (2020) an und differenziert diese für die spezifischen Anforderungen des Orthographieerwerbs weiter. Orthographische Abweichungen werden in DeChiLko in fünf Bereiche sowie 20 Annotationsebenen unterteilt. Diese Multiebenen-Annotation erlaubt eine detaillierte Erfassung jedes Rechtschreibfehlers, auch wenn es dabei zu Überlappungen zwischen den Kategorien kommen kann. Dieses Schema ermöglicht eine differenzierte Analyse der

orthographischen Kompetenz chinesischer Lernender und leistet damit einen bedeutenden Beitrag zur Erforschung ihrer Rechtschreibleistungen.

Tabelle 15 stellt einen Überblick über das Annotationsschema des DeChiLKo-Korpus dar. Mithilfe des daraus entwickelten Tagsets werden die orthographischen Abweichungen im zweiten Schritt der Annotation erfasst.

Tabelle 15: Überblick über das Annotationsschema des DeChiLKo-Korpus

Schicht	Annotationsebene	Erläuterung
tok	word	Originaltext mit aufeinanderfolgenden Tokens
	S	Satzspannen
	pos	Originaltext: Treetagger-POS-Tags (STTS)
	lemma	Originaltext: Treetagger-Lemmata
ZH	ZH	ZH: Zielhypothese
	ZHDiff	ZH: Abweichungen ZH – word
	ZHS	ZH: Satzspannen
	ZHpos	ZH: Treetagger-POS-Tags (STTS)
	ZHlemma	ZH: Treetagger-Lemmata
Phonographischer Bereich	Graph_Fehler	Segmentierung des Tokens nach Graphem oder Graphemfolgen.
	Graph_ZH	Segmentierung der ZH nach Graphem oder Graphemfolgen
	Graph_PGK	Phonem-Graphem-Korrespondenz
	Graph_BF	Buchstabenform
	Graph_VQM	Vokalquantitätsmarkierung
	Graph_S-Schreibung	s-Schreibung
	Graph_SG	Spezielle Grapheme
Silbischer Bereich	Silben_Fehler	Segmentierung des Tokens nach Silbenstruktur
	Silben_ZH	Segmentierung der ZH nach Silbenstruktur
	Silben_AR	Anfangsrand
	Silben_SK	Silbenkern
	Silben_ER	Endrand

(Fortsetzung Tabelle 15)

Schicht	Annotationsebene	Erläuterung
Morphologischer Bereich	Morph_Fehler	Segmentierung des Tokens nach Morphemen
	Morph_ZH	Segmentierung der ZH nach Morphemen
	Morphem	Abweichungen nach Morphemarten ¹¹
	Morph_Konstanz	Morphologische Konstanzschreibung
	Morph_KompoS	Kompositumschreibung
	Morph_AffixS	Affixschreibung
Syntaktischer Bereich	Syn_GKS	Groß- und Kleinschreibung
	Syn_GZS	Getrennt- und Zusammenschreibung
	Syn_dass-das	Verwechslung zwischen <i>dass</i> und <i>das</i>
	Syn_man-Mann	Verwechslung zwischen <i>man</i> und <i>Mann</i>
	Syn_SZ	Satzzeichen
Sonstiges	Son_FW	Fremdwortschreibung
	Son_Gra-Ab	Grammatisch abzuleitende Schreibung wie <i>dem</i> statt <i>den</i>
	SON	Sonstige Schreibung

Mit diesem Annotationsschema wird beispielsweise die Schreibung *<farat> auf mehreren Ebenen analysiert und annotiert:

1. Phonographischer Bereich:

- VR-: fehlendes vokalisiertes <r> auf der [Graph_PGK]-Ebene: *<a> statt <ahr>
- VM-: fehlende Markierung des Langvokals auf der [Graph_VQM]-Ebene: *<a> statt <ah>

2. Silbischer Bereich:

- *ZDiph: fehlerhafte Markierung von zentralisierendem Diphthong im Silbenkern ([Silben_SK]-Ebene): *<a> für <ahr>
- *Ein: fehlerhafter einfacher Endrand auf der [Silben_ER]-Ebene: *<t> für <d>

3. Morphologischer Bereich:

- *LM: falsches lexikalisches Morphem auf der [Morphem]-Ebene: *{fa} für {fahr} und *{rad} für {rat}

¹¹ Falls es sich bei den Abweichungen nicht um spezifische morphologische Schreibungen handelt, werden sie in der Annotationsebene [Morphem] lediglich nach Morphemarten als Abweichungen bei den lexikalischen Morphemen (LM) oder grammatischen Morphemen (GM) zugeordnet.

- DH: Fehler bei der morphologischen Dehnungsschreibung auf der Ebene [Morph_Konstanz]: *{fa} statt {fahr}
- ALV: Fehler bei der Auslautverhärtung auf der Ebene [Morph_Konstanz]: *{rat}für {rad}
- PV: Konsonantenauslassung bei Phonemverschmelzung in der Kompositumschreibung [Morph_KompoS]: **farat* für *Fahrrad*

4. Syntaktischer Bereich:

- KfG_Sub: Kleinschreibung für Großschreibung bei Substantiven auf der Ebene [Syn_GKS]

In EXMARaLDA (Dulko) sehen die Annotationen des Beispielausschnitts wie der Screenshot in Abbildung 8 aus:

[Graph_Fehler]	f	a	ren	f	a		ra	t		ko	h	en	
[Graph_ZH]	f	ah	ren	F	ah	r	ra	d		ko	ch	en	
[Graph_PGK]						VR-					*K		
[Graph_BF]													
[Graph_KLM]			VM-			VM-							
[Graph_S-Schreibung]													
[Graph_SG]													
[Silben_Fehler]	f	a	ren	f	a		ra	t		ko	h	en	
[Silben_ZH]	f	ah	ren	F	ahr		ra	d		ko	ch	en	
[Silben_AR]													
[Silben_SK]			*Mono			*ZDiph							
[Silben_ER]								*ein			SIG		
[Morph_Fehler]	far		en	fa			rat			koh		en	
[Morph_ZH]	fahr		en	Fahr			rad			koch		en	
[Morphem]	*LM			*LM			*LM			*LM			
[Morph_Konstanz]				DH			ALV						
[Morph_KompoS]				PV									
[Morph_AffixS]													
[Syn_GKS]				KfG-Sub									
[Syn_GZS]													
[Syn_dass-das]													
[Syn_man-Mann]													
[Syn_SZ]													
[Son_Gra-Ab]													
[Son_FW]													
[SON]													

Abbildung 8: Screenshot der orthographischen Annotationen des Lernertexts (C_NW_2019_D_007) mit den Spuren wie [Graphem], [Graph_BF], [Graph_VQM] usw., auf denen die orthographischen Abweichungen feiner annotiert werden. Einige der zuvor benannten Annotationsebenen wurden hier zugunsten der besseren Übersicht weggelassen.

3.3 Korpusabfragen in ANNIS

Nach der Transkription und Annotation lassen sich die Lernertexte in EXMARaLDA (Dulko) in einem spezifischen Format *.exb* speichern. Um das gesamte DeChiLKo-Kor-

pus später für die Korpusnutzung auf der ANNIS-Suchplattform (Version 3.7.1) (Krause & Zeldes, 2016) verfügbar zu machen, werden alle annotierten Lernertexte durch die Anwendung des von Andreas Nolda entwickelten Build-Systems *Makedulko*¹² in das ANNIS-Format konvertiert.

ANNIS (ANNotation of Information Structure) ist ein webbasiertes Werkzeug für die Durchsuchung und Visualisierung linguistischer Korpora. Mit ANNIS können Nutzer*innen anspruchsvolle Suchabfragen in mehrschichtigen Korpora durchführen und die Ergebnisse in diversen visuellen Darstellungen aufbereiten. Ein signifikanter Vorteil von ANNIS ist die Funktion, die Durchführung kombinierter Korpusuchen über verschiedene Ebenen hinweg zu ermöglichen (Krause & Zeldes, 2016).

In Tabelle 16 sind exemplarische Suchszenarien dargestellt.

Tabelle 16: Exemplarische Suchanfragen

Sucheingabe	Bemerkung	Beispieltreffer
word=„farat“	Exakte Wortform in den Lernertexten	<i>farat</i>
lemma=„fahren“	Lemmbasierte Suche in den Lernertexten	<i>fährt, fährt, fahren, Fahren...</i>
ZH=„Fahrrad“	Exakte Wortform in der Zielhypothese	<i>Fahrrad</i> ¹³
ZHlemma=„eine“	Lemmbasierte Suche auf der Zielhypotheseebene	<i>ein, Ein, eine, Eine, einen, Einen...</i>
pos=„KOUS“	Abfrage nach Wortart	<i>als, dass, weil...</i>
ZHpos=„NN“	Abfrage nach Nomen auf der Zielhypotheseebene	<i>Sonntag, Teehaus...</i>
ZHDiff=„MERG“	Suche nach Änderungen auf der [ZHDiff]-Ebene	<i>einer seits statt einerseits</i>
Graphem=„K+“	Suche nach überflüssigem Konsonantengraphem bzw. überflüssigen Konsonantengraphemen auf der [Graphem]-Ebene	Überflüssiges <G> in *<Grund> für <i>Rund</i>
Graph_VQM=“*KD“	Suche nach fehlerhafter Konsonantenverdopplung auf der Kürze-Längenmarkierungsebene	<ll> für <l> in *<schmalen> für <i>schmalen</i>
Graph_SG =„FVS“ _=_Silben_AR	Verwechslung zwischen <f> und <v> am Silbenanfangsrand	<v> für <f> in *<versetzen> für <i>fortsetzen</i>
Morph_AffixS=“*Prä” _i_ Graph_SG =„FVS“	Verwechslung zwischen <f> und <v> bei einer falschen Präfixschreibung	<f> für <v> in *<foreisammen> für <i>vereinsamen</i>

¹² Mehr Information unter: <https://sr.ht/~nolda/makedulko/> (13.06.2026)

¹³ Eine Suchanfrage auf der [ZH]-Ebene bezieht sich auf die exakte Wortform in der Zielhypothese selbst. Wird auf der [ZH]-Ebene nach ZH=„Fahrrad“ gesucht, so wird der exakte Eintrag „Fahrrad“ als Treffer gefunden. Dies ermöglicht es, gleichzeitig die zugehörigen Originalformen auf der [Wort]-Ebene, die dieser Zielhypothese zugeordnet sind – wie beispielsweise *<farat>, *<fahrad>, *<faahrt> usw. und die korrekte Form *Fahrrad* – indirekt mitzufinden.

(Fortsetzung Tabelle 16)

Sucheingabe	Bemerkung	Beispieltreffer
Syn_GKS=„GfK“ _= _ZHpos=/V.*/	Großschreibung für Kleinschreibung bei den Verben	<i>er *Plant einen Ausflug; Sie werden uns *Besuchen</i>
Son_Gra-Ab=„Dek“ _i_ Graph_Fehler=/(m n)/ _= _Graph_ZH=/(m n)/	Verwechslung zwischen den Graphemen <m> und <n> wegen eines Deklinationsfehlers	<i>*dem statt den; *Im statt In</i>
Syn_dass-das != „“ & meta::institution_category = „B“	Verwechslung zwischen <i>dass</i> und <i>das</i> in den Lernertexten aus Hochschulen der Kategorie B.	<i>[...], *das Kinder vor dem Computer vereinsamen.</i>

Die Suchanfragen in ANNIS können durch die Auswahl spezifischer Korpora entweder im gesamten DeChiLko oder in beliebigen Subkorpora durchgeführt werden. Abbildung 9 zeigt beispielsweise die Suchergebnisse zur Groß- statt Kleinschreibung in Lernertexten aus Hochschulen der Kategorie A in den ausgewählten Subkorpora – dem Erwerbskorpus und dem Prüfungskorpus_2019.

word _= _Syn_GKS="GfK"
& meta::institution_category =
"A"

Suchanfrage

Q Search More History

90 matches
in 55 documents

Corpus List Search Options

Visible: All

Filter

Name	Texts	Tokens
DeChiLko_2019-v1.0	95	12.786
DeChiLko_Erwerb-v1.0	140	17.480
DeChiLko_Prüfung2017-v1.0	100	11.540

Auswahl der Korpora zur
Durchsuchung

Help/Examples Query Builder Query Result

Base text

Displaying Results 1 - 10 of 90

1 Path: DeChiLko_2019-v1.0 > A_O_2019_D_001 (tokens 2 - 12)

word % der Bundesbürger sehen regelmäßig Fern Sie wollen sich am

grid (default_ns)

pos	NN	ART	NN	VVFIN	ADJD	VVFIN	\$	PPER	VMFIN	PRF	APPRART
lemma	%	die	Bundesbürger	sehen	regelmäßig	Fern	.	Sie	wollen	sich	an+die
S	s1							s2			
ZH	%	der	Bundesbürger	sehen	regelmäßig	fern	.	Sie	wollen	sich	am
ZHpos	NN	ART	NN	VVFIN	ADJD	PTKVZ	\$	PPER	VMFIN	PRF	APPRART
ZHlemma	%	die	Bundesbürger	sehen	regelmäßig	fern	.	Sie	wollen	sich	an+die
ZHs	s1							s2			
ZHdiff								CHA			
Syn_GKS								GfK			

2 Path: DeChiLko_2019-v1.0 > A_O_2019_D_002 (tokens 84 - 94)

word Seits müssen die Leute aber Sparen Immer mehr Deutsche gehen

grid (default_ns)

3 Path: DeChiLko_2019-v1.0 > A_O_2019_D_003 (tokens 2 - 12)

word Prozent der Bundesbürger sehen regelmäßig Fern Sie wollen sich am

grid (default_ns)

4 Path: DeChiLko_2019-v1.0 > A_O_2019_D_009 (tokens 2 - 12)

word % der Bundesbürger sehen regelmäßig Fern Sie wollen sich am

grid (default_ns)

5 Path: DeChiLko_2019-v1.0 > A_ZS_2019_D_003 (tokens 106 - 112)

word Freizeit

grid (default_ns)

Suchergebnisse zur Groß- statt Kleinschreibung in Lernertexten
aus Hochschulen der Kategorie A in den ausgewählten Korpora
(Erwerbskorpus und Prüfungskorpus 2019)

Abbildung 9: Screenshot der Abfrage nach Groß- für Kleinschreibungen in den Lernertexten aus Hochschulen der Kategorie A in den ausgewählten Korpora (Erwerbskorpus und Prüfungskorpus_2019) im ANNIS-Interface

Nach der Konvention mit *MakeDulko* besteht das gesamte DeChiLko in ANNIS aus 329 Diktattexten; die sämtlichen Tokens betragen 31.674. Das Prüfungskorpus umfasst 195 Lernertexte, davon 100 Diktate aus dem Jahr 2017 und 95 aus dem Jahr 2019. Das Erwerbskorpus besteht aus 140 Lernertexten, die kontinuierlich über drei Semester

(Wintersemester 2021/22 bis zum Wintersemester 2022/23) erhoben wurden. Die Anzahl der Tokens in den 195 Texten im Prüfungskorpus beträgt 17.598, während die Anzahl der Tokens im Erwerbskorpus 14.076 beträgt. Die weiteren Daten über die zwei Subkorpora werden in der folgenden Tabelle 17 dargestellt.

Tabelle 17: Überblick über das DeChiLko-Korpus

	Lernertexte (N Probanden)		Σ	Tokens		Σ
	2017	2019		2017	2019	
Prüfungskorpus	100	95	195	8.219	9.379	17.598
Erwerbskorpus	140			14.076		
Σ	329			31.674		

3.4 Methode zur Analyse der Korpusdaten

Mit der Korpusabfrage nach Wörtern im Lernertext, Fehlerannotationen oder Kombinationen von mehreren Annotationsebenen sowie Metadaten in ANNIS wird die orthographische Leistung chinesischer Deutschlernender quantitativ und qualitativ analysiert. Die Analyse gliedert sich in zwei Hauptteile. Der erste Teil der Analyse ist eine deskriptive Auswertung, die die Häufigkeit und Verteilung der orthographischen Fehler in den verschiedenen linguistischen Bereichen darstellt. Als zentraler deskriptiv-statistischer Kennwert dient hierbei der Fehlerwert pro 1.000 Tokens dazu, die Fehlerhäufigkeit über Texte unterschiedlicher Länge hinweg vergleichbar zu machen. Ferner werden die typischen Fehlerschreibungen einzelner Kategorien anhand von Beispielen aus dem Korpus qualitativ dargestellt und analysiert.

Der zweite Teil besteht aus einer inferenzstatistischen Überprüfung der Korrelation zwischen der Orthographiekompetenz und verschiedenen Faktoren. Für diese Analyse wurde aufgrund der hohen Anzahl an Nullwerten in den Fehlerdaten ein zweistufiges Hurdle-Modell gewählt.

- **Stufe 1** modelliert mittels eines generalisierten linearen gemischten Modells (GLMM) die Wahrscheinlichkeit, ob ein bestimmter Fehler überhaupt auftritt (Fehlerwert > 0).
- **Stufe 2** analysiert mittels eines linearen gemischten Modells (LMM) die Höhe des Fehlerwertes nur für jene Fälle, in denen tatsächlich ein Fehler vorliegt.

Als Einflussfaktoren (feste Effekte) wurden im Prüfungskorpus die Variablen *Diktattext*, *Prüfungsstatus*, *Hochschulkategorie* und *Region* berücksichtigt. Im Erwerbskorpus waren dies die Faktoren *Lerndauer* und *Aufgabentyp*. Die verschiedenen Fehlerkategorien sowie die einzelnen Lernenden (im Erwerbskorpus) wurden als zufällige Effekte modelliert, um deren Varianz zu kontrollieren. Die Signifikanz der einzelnen Faktoren

wurde mittels Likelihood-Quotienten-Tests ermittelt. Die statistische Analyse und die grafische Darstellung der Daten erfolgten in R¹⁴ (R Core Team, 2022) und R-Studio¹⁵ (Posit team, 2023) unter Zuhilfenahme des Pakets *lme4* (Bates et al., 2015) für die gemischten Modelle.

¹⁴ Weitere Information zu R unter: <https://www.r-project.org/> (13.06.2026)

¹⁵ Weitere Information zu R-Studio unter: <http://www.posit.co/> (13.06.2026)

4 Orthographische Fehler der chinesischen Deutschlernenden

In diesem Kapitel wird zunächst ein Überblick über die allgemeine Verteilung der orthographischen Fehler im Korpus geschildert. Anschließend werden einzelne Rechtschreibphänomene tiefgehend in fünf orthographischen Teilbereichen analysiert.

4.1 Überblick

Die [ZHDiff]-Ebene ist eine der wichtigsten Annotationsebenen, da sie Abweichungen zwischen den von den Lernenden verfassten Diktaten und dem Diktattext (Zielhypothese) auf Tokenebene erfasst. Die Verteilung der Annotationen auf der [ZHDiff]-Ebene gibt einen ersten Überblick über die orthographische Leistung: So lässt sich die Annotation „CHA (change)“ mit hoher Wahrscheinlichkeit einem Schreibfehler (im Bereich Wortorthographie, Zeichensetzung oder Groß- und Kleinschreibung) zuordnen. Die Annotationen „SPLIT“ und „MERG“ signalisieren typischerweise Fehler im Bereich der Getrennt- und Zusammenschreibung.

Insgesamt werden 6885 Treffer in der [ZHDiff]-Ebene von 333 Texten in dem gesamten DeChiLko-Korpus mit 31.674 Tokens gefunden. Der durchschnittliche Fehlerwert beträgt 217,37 pro 1.000 Tokens. Jedoch darf man nicht ignorieren, dass in zwei Lernertexten (Prüfungskorpus_A_ZS_2017_002, Erwerbskorpus_LWY_20211217) überhaupt keine Abweichungen zu finden sind.

Tabelle 18: Überblick über [ZHDiff]-Annotationen

[ZHDiff]-Annotationen insgesamt	Summe aller Tokens	[ZHDiff]-Annotationen pro 1.000 Tokens
6.885	31.674	217,37

Von dieser Statistik kann man deutlich erkennen, dass die chinesischen Deutschlernenden im Allgemeinen Schwierigkeiten in der Orthographie haben. Die Verteilung der [ZHDiff]-Annotationen pro 1.000 Tokens wird im folgenden Säulendiagramm dargestellt.

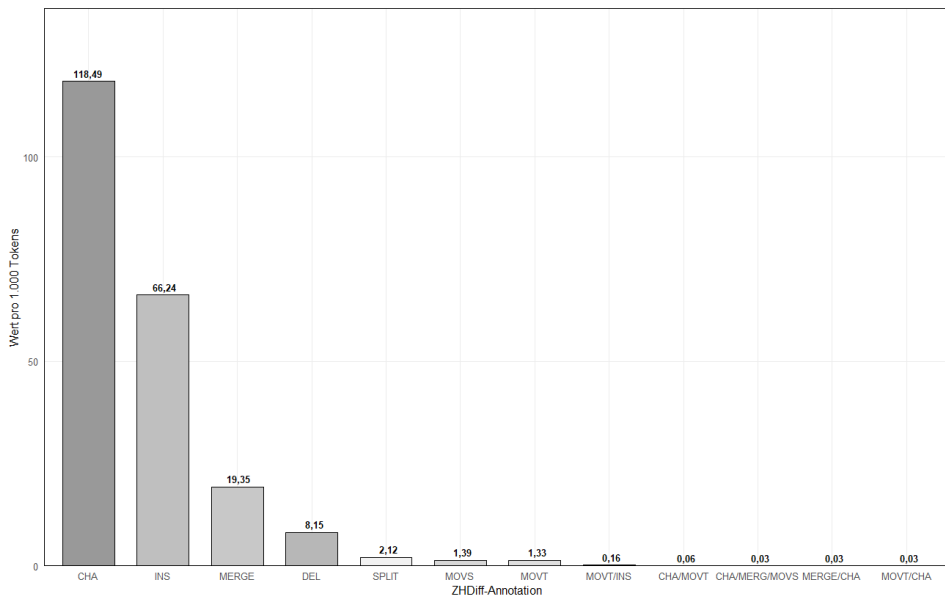


Abbildung 10: Verteilung der [ZHDiff]-Annotationen; Wert pro 1.000 Tokens

Die Kategorie „CHA“ (Change) weist mit einem Fehlerwert von 118,49 pro 1.000 Tokens den höchsten Wert unter allen Annotationen auf. Es deutet darauf hin, dass die Lernenden Änderungen vornehmen müssen, um ihre Schreibung an die Schreibnorm anzupassen. Darauf folgt die Kategorie „MERG“ (Merg) mit einem Wert von 19,35 pro 1.000 Tokens, während die Fehlerkategorie „SPLIT“ mit einem geringeren Fehlerwert von 2,12 pro 1.000 Tokens ebenfalls vertreten ist. Diese Daten deuten auf spezifische Herausforderungen der Lernenden im Bereich der Getrennt- und Zusammenschreibung hin.

Der hohe Fehlerwert von 66,24 pro 1.000 Tokens in den Kategorien „INS“ (Insert) und „DEL“ (deletion) zeigt, dass die Lernenden während des Schreibens ein bestimmtes Token unnötigerweise hinzufügen oder auslassen. Da sich diese Arbeit auf orthographische Abweichungen konzentriert, werden bei der Annotation insbesondere ausgelassene und überflüssige Interpunktionen im Hinblick auf die Zeichensetzung betrachtet.

Die Analyse der Wortartverteilung auf der Annotationsebene [ZHpos] (siehe Abbildung 11) belegt, dass Abweichungen der Kategorie „CHA“ am häufigsten bei Nomen und Verben vorkommen. Dieser Befund korreliert mit der Verteilung der Inhaltswörter im Deutschen: Unter den lexikalischen Wortklassen zählen Substantive und Verben zu den wortschatzstärksten Kategorien (vgl. Duden, o. J.), weshalb auch die absolute Anzahl orthographischer Abweichungen in diesen Klassen erwartungsgemäß am höchsten ausfällt. Funktionswörter wie Pronomen, obwohl tokenhäufig, weisen aufgrund ihrer begrenzten orthographischen Varianz deutlich weniger Abweichungen auf. Die Kategorie „MERG“ tritt überwiegend bei Nomen auf. Eine qualitative Sichtung

der betroffenen Fälle zeigt, dass es sich dabei überwiegend um Nominalkomposita handelt, die fälschlicherweise getrennt geschrieben wurden (z. B. **Haus Tier* statt *Haustier*). Da Nominalkomposita im Deutschen eine hochproduktive Wortbildungsart darstellen und eine verbindliche Zusammenschreibung erfordern, ist ihre Überrepräsentation in der Fehlerkategorie MERG plausibel.

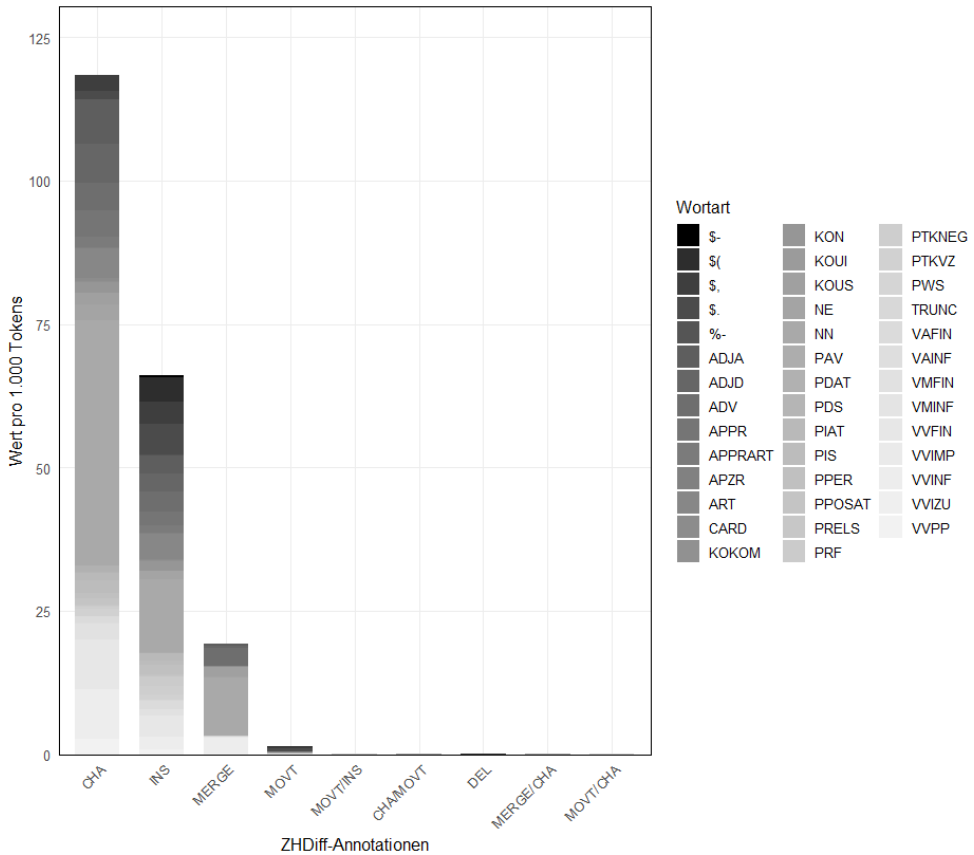


Abbildung 11: Wortartverteilung zu den [ZHDiff]-Annotationen; Wert pro 1.000 Tokens

Die Analyse der [ZHDiff]-Annotationen ermöglicht zwar einen Einblick in die Abweichungen der Lernerproduktionen auf Token-Ebene im Vergleich zur Zielhypothese, jedoch wird hier das Token als eine umfassende Einheit betrachtet. Um ein vollständiges Bild der orthographischen Kompetenz chinesischer Deutschlerner zu erhalten, ist es erforderlich, die Analyse über die Token-Ebene hinaus zu erweitern, weitere orthographische Annotationsebenen zu berücksichtigen und eine detaillierte Analyse der orthographischen Abweichungen an jeder Fehlerstelle vorzunehmen. Die folgende Abbildung 12 veranschaulicht die Fehlerverteilung in den einzelnen orthographischen Bereichen.

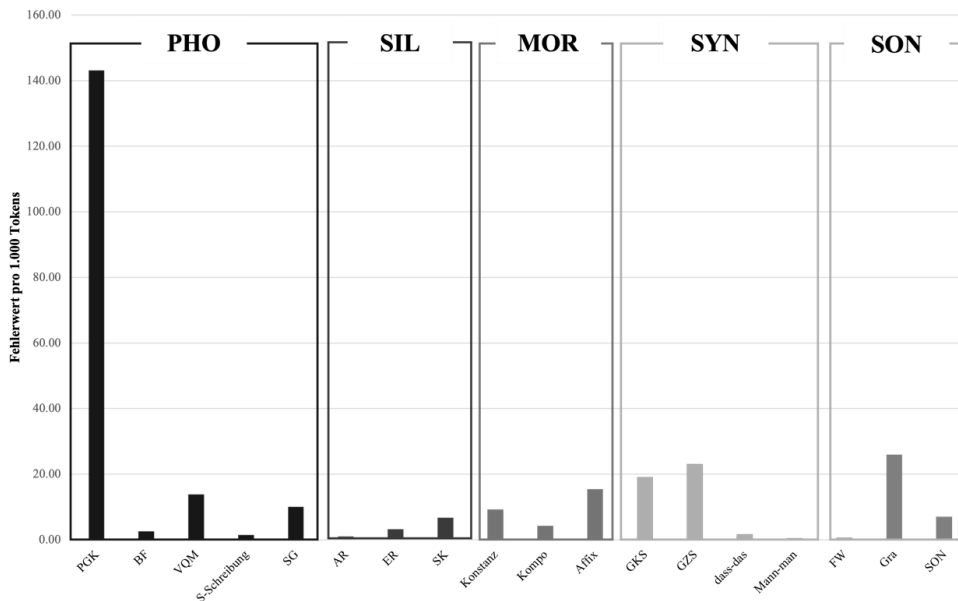


Abbildung 12: Überblick über Fehlerverteilungen nach Fehlerkategorien in den orthographischen Bereichen [PHO], [SIL], [MOR], [SYN], [SON]; Fehlerwert pro 1.000 Tokens

Die folgende Detailanalyse der orthographischen Fehler basiert auf zwei methodischen Vorbemerkungen. Zum einen können die Annotationsebenen Überlappungen aufweisen (siehe Kapitel 3.4, Abbildung 8). Zum anderen werden alle Fehlerfrequenzen auf 1.000 Tokens normiert, um eine valide Vergleichbarkeit über die unterschiedlich langen Lernertexte hinweg zu gewährleisten.

4.2 Der phonographische Bereich

Der phonographische Bereich der deutschen Orthographie umfasst Regeln, die primär auf der Beziehung zwischen Phonem und Graphem basieren. Abweichungen in diesem Bereich deuten vor allem darauf hin, dass die Lernenden Schwierigkeiten haben, die gesprochene Sprache adäquat in die Schrift umzusetzen oder spezifische graphematische Konventionen korrekt anzuwenden, die eng mit der Lautstruktur verbunden sind.

In diesem Kapitel werden die Fehler chinesischer Deutschlernender im phonographischen Bereich detailliert analysiert. Die Untersuchung stützt sich auf mehrere

Annotationsebenen, die im DeChiLKo-Korpus zur Erfassung spezifischer orthographischer Phänomene etabliert wurden:

- [Graph_PGK] (Phonem-Graphem-Korrespondenz): Abweichungen von den grundlegenden Laut-Buchstaben-Zuordnungen;
- [Graph_BF] (Buchstabenform): Fehler bei der korrekten graphischen Realisierung einzelner Buchstaben, einschließlich der Setzung von Umlautpunkten;
- [Graph_VQM] (Vokalquantitätsmarkierung): Fehler bei der orthographischen Kennzeichnung von Vokallänge und -kürze;
- [Graph_S-Schreibung]: spezifische Fehler bei der komplexen Schreibung der s-Laute (<s>, <ss>, <ß>);
- [Graph_SG] (Spezielle Grapheme): Fehler bei der Schreibung von Graphemen mit besonderen oder weniger transparenten Phonem-Graphem-Beziehungen (z. B. <v>/<f>, <e>/<ä>).

Die Ebenen [Graph_Fehler] und [Graph_ZH] dienen dabei der grundlegenden Segmentierung des Lernertexts bzw. der Zielhypothese nach Graphemen oder Graphemfolgen und bilden die Basis für die detailliertere Fehlerklassifikation auf den obengenannten Ebenen. Ziel der folgenden Abschnitte ist es, die spezifischen orthographischen Phänomene, die auf diesen Annotationsebenen beobachtet wurden, quantitativ und qualitativ zu analysieren.

4.2.1 Fehler der Buchstabenform [Graph_BF]

Die Annotationsebene [Graph_BF] erfasst Abweichungen von der normgerechten graphemischen Form der Buchstaben des deutschen Alphabets. Hierzu zählen fehlerhafte Realisierungen wie spiegelbildliche oder unvollständige Buchstaben (annotiert als „*BF“) sowie Abweichungen bei der Setzung von Umlautpunkten (Auslassung oder Hinzufügung, annotiert als „UP“). Im gesamten DeChiLKo-Korpus werden 79 Treffer in 68 verschiedenen Lernertexten identifiziert. Dies entspricht einem sehr geringen Fehlerwert von 2,53 pro 1.000 Tokens. Die überwiegende Mehrheit dieser Fehler, nämlich 98,7 %, betraf die Auslassung oder fälschliche Hinzufügung von Umlautpunkten.

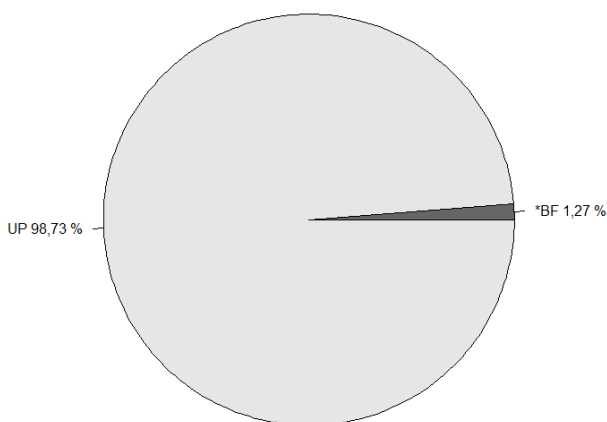


Abbildung 13: Prozentuale Fehlerverteilung in der Kategorie Buchstabenform

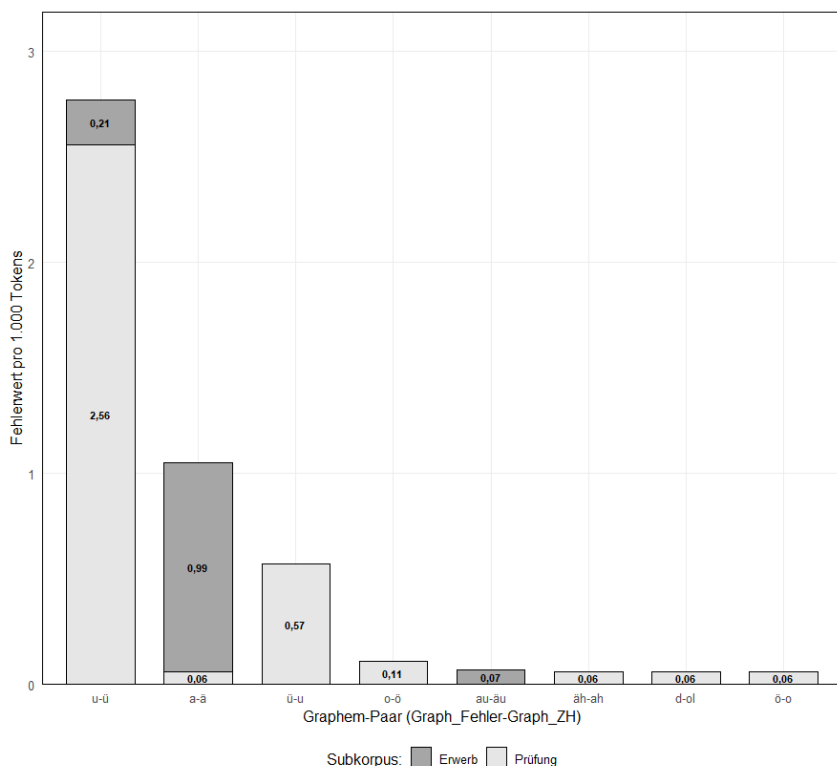


Abbildung 14: BF-Fehlerverteilung in den DeChiLKo-Subkorpora¹⁶; Fehlerwerte pro 1.000 Tokens

¹⁶ Anmerkung zu den Abbildungen: Für alle Grafiken, die einen Vergleich zwischen den Subkorpora darstellen, ist zu beachten, dass die Subkorpora in Bezug auf ihre Größe variieren. Um eine vergleichbare Darstellung der Fehlerhäufigkeiten über verschiedene Subkorpora hinweg zu gewährleisten, werden die Fehlerwerte und andere statistische Maße in Relation zur Gesamtzahl der Tokens im jeweiligen Subkorpus berechnet.

Die Auslassung des Umlautpunktes bei <ü> trat dabei am häufigsten auf und war in allen Subkorpora zu beobachten. Danach folgt die Auslassung des Umlautpunkts bei <ä>, und zwar vorwiegend im Erwerbskorpus. Auch die überflüssige Hinzufügung von Umlautpunkten bei <u> wurde vereinzelt festgestellt (siehe Abbildung 14). Im Prüfungskorpus (D_ZS_2017_001) wurde eine einzige Abweichung der Kategorie „*BF“ (fehlerhafte Buchstabenform) identifiziert: Durch das Auslassen des Zwischenraums verschmolz die Buchstabenkombination visuell zum Graphem <d>¹⁷.

Die geringe Fehlerzahl im Bereich der grundlegenden Buchstabenform korrespondiert mit Befunden aus Studien zu L1- und L2-Lernern des Deutschen (vgl. Thomé, 1987: 102 f.). Für die untersuchten chinesischen Studierenden, deren L1 Chinesisch keine Alphabetschrift ist, ist dieses Ergebnis bemerkenswert. Es deutet darauf hin, dass die grundlegende Beherrschung lateinischer Buchstabenformen bereits stark gefestigt ist. Es ist anzunehmen, dass das graphetische Wissen, das die Lernenden durch den Erwerb des *Pinyin*-Systems und durch den Englischunterricht aufgebaut haben, auf das Deutsche als L3 übertragen wird (vgl. Fan, 2017; Lü, 2017). Hierbei sorgt insbesondere das *Pinyin*-System für die Grundlage des alphabetischen Prinzips (vgl. Lü, 2017; Zhou & McBride, 2022), während das Englische als typologisch nahe L2 als Brücke fungiert (vgl. Fan, 2017; Hufeisen & Neuner, 2005). Dies deutet darauf hin, dass die „Alphabetisierungsphase“ (im Sinne der Beherrschung des lateinischen Alphabets) bereits vor dem Deutschlernen abgeschlossen ist. Das L1-Wissen (*Pinyin*) und das L2-Wissen (Englisch) bilden hier eine solide Basis für die graphetische Formkorrektheit im Deutschen.

Die deutliche Dominanz von Fehlern bei der Umlautkennzeichnung (97,5 %) lenkt den Fokus auf ein spezifisch deutsches Orthographiephänomen. Umlaute und ihre graphemische Markierung durch Punkte über den Vokalbuchstaben <ä, ö, ü> stellen eine Besonderheit der deutschen Schrift dar, die in den Herkunftssprachen der untersuchten Lernenden – Chinesisch (L1) sowie *Pinyin* und Englisch (L2) – keine funktionalen Entsprechungen hat. Zwar existiert im *Pinyin* <ü>, jedoch unterscheiden sich sowohl die phonologische und graphematische Funktion grundlegend von der deutschen Umlautschreibung. Der Einfluss der L1 und L2 ist hier also weniger als direkte Interferenzquelle für die Form der Umlautpunkte zu sehen, sondern eher im Fehlen entsprechender Konzepte, was die Aufmerksamkeit für dieses spezifisch deutsche Merkmal möglicherweise nicht von vornherein schärft.

Die weiterführende Analyse im DeChiLKo-Korpus ergab, dass 85,9 % aller Abweichungen bei Umlautpunkten mit der korrekten Markierung morphologischer Umlautungen (n = 67, erfasst auf der Ebene [Morph_Konstanz]) zusammenhängen. Über ein Viertel der Umlautfehler (n = 20) ließ sich zudem direkt auf Konjugations- (n = 13) oder Deklinationsfehler (n = 7) (erfasst auf der Ebene [Son_Gra-Ab]) zurückführen. Dies stützt die Annahme, dass die Schwierigkeit weniger in der graphemischen Realisierung der Umlautpunkte liegt, sondern vielmehr in der korrekten Anwendung der zugrunde liegenden deutschen morphologischen und grammatischen Regeln, die eine Umlautung erfordern. Die Fehlerquelle liegt somit weniger auf der graphematischen

17 Bei dieser Abweichung könnte es sich um ein Schreibversehen handeln. Um jedoch die Vollständigkeit der Fehlererfassung zu gewährleisten, wurde sie im Korpus der Kategorie „fehlerhafte Buchstabenform“ (*BF) zugeordnet.

Ebene der Umlautmarkierung selbst – also der Frage, wie Umlaute zu schreiben sind, sondern vielmehr auf der morphologisch-grammatischen Ebene: Die Lernenden wissen offenbar, dass Umlaute durch Umlautpunkte markiert werden, wenden dieses Wissen jedoch nicht korrekt an, weil das zugrunde liegende Regelwissen über umlautauslösende Prozesse im Deutschen noch nicht hinreichend internalisiert ist. Dies verdeutlicht die enge Interaktion verschiedener sprachlicher Ebenen (Orthographie, Morphologie, Grammatik) innerhalb des L2-Erwerbsprozesses, wie sie auch im Rahmen des Multi-Kompetenz-Ansatzes postuliert wird, wo verschiedene Wissensmodule interagieren. Die Lernenden müssen nicht nur die Form der Umlaute kennen, sondern vor allem das komplexe Regelwissen des Deutschen internalisieren.

4.2.2 Fehler bei der Phonem-Graphem-Korrespondenz [Graph_PGK]

Die Annotationsebene [Graph_PGK] erfasst Abweichungen von der grundlegenden Phonem-Graphem-Korrespondenz (PGK-Regeln) des Deutschen. Im Fokus stehen hier die direkten, primären Laut-Buchstaben-Beziehungen (Drosdowski & Eisenberg, 1995), d.h. die lautgetreue Verschriftung. Komplexere Schreibweisen, die übergeordneten Prinzipien folgen (z. B. spezifische Markierungen der Vokalquantität oder die *s*-Schreibung), werden in den nachfolgenden Abschnitten (4.2.3–4.2.5) gesondert behandelt. Die im Rahmen von [Graph_PGK] erfassten Abweichungen wurden nach Fehlerart (fehlerhaft „*“, überflüssig „+“, ausgelassen „-“) und betroffener Lautkategorie (Konsonant K, Vokal V, Reduktionsvokal RV, vokalisiertes <r> VR, Graphemvertauschung GV) klassifiziert.

Mit insgesamt 4.529 Abweichungen in 326 Texten erweist sich die grundlegende Phonem-Graphem-Korrespondenz als eine zentrale Herausforderung für chinesische Deutschlernende. Der Fehlerwert von 142,99 pro 1.000 Tokens unterstreicht die hohe Fehleranfälligkeit in diesem Aspekt der deutschen Orthographie.

Tabelle 19: Fehlerwert der PGK-Abweichungen pro 1.000 Tokens

Graphem-Annotationen insgesamt	Summe aller Tokens	Graphem-Annotationen pro 1.000 Tokens
4.529	31.674	142,99

Die Verteilung der PGK-Abweichungen nach Teilkategorien in Abbildung 15 zeigt, dass Abweichungen bei Konsonanten mit einem Wert von 65,92 pro 1.000 Tokens die mit Abstand häufigste Fehlerkategorie darstellen. Darauf folgen Abweichungen bei der Verschriftlichung von Reduktionsvokalen (39,44) und Vokalen (27,01). Der Fehlerwert beim vokalisiertem <r> beträgt 9,65 pro 1.000 Tokens, während Graphemvertauschungen (1,04) eine marginale Rolle spielen.

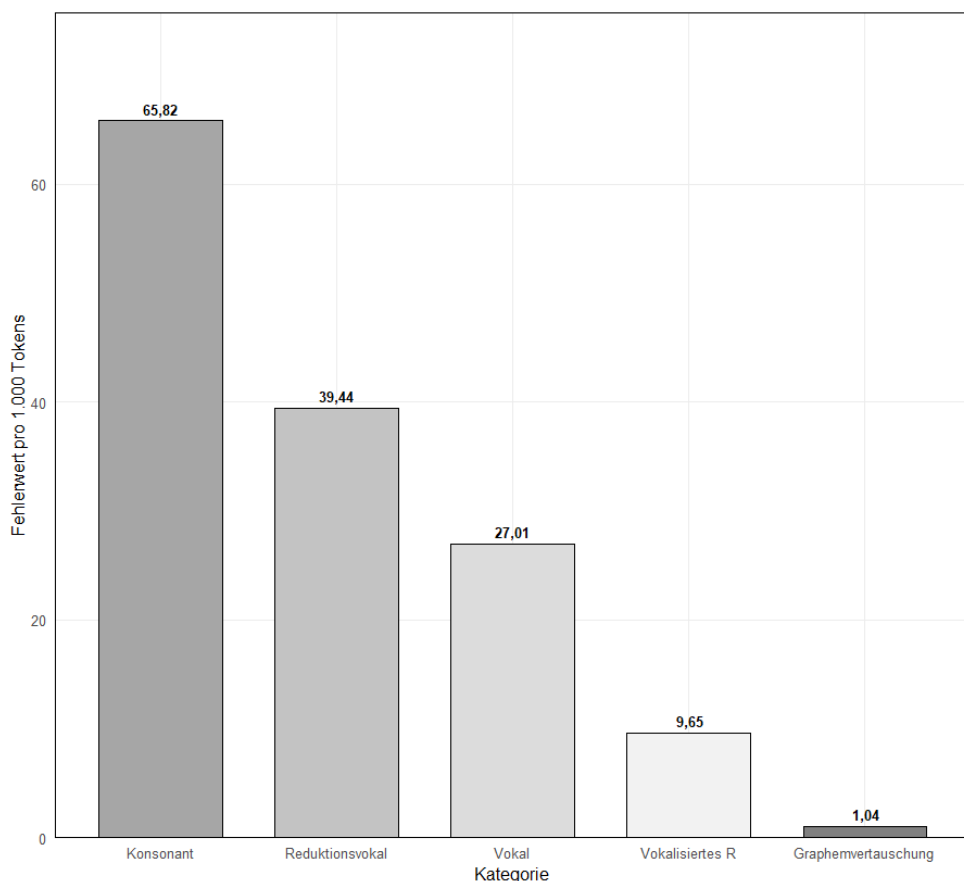


Abbildung 15: PGK-Fehlerverteilung nach Teilkategorien; Fehlerwert pro 1.000 Tokens

Eine detailliertere Betrachtung der Fehlerarten in Abbildung 16 verdeutlicht, dass fehlende Konsonantengrapheme (K-) die häufigste Fehlerart sind. Ebenfalls bemerkenswert häufig sind die fehlerhafte Verschriftlichung von Konsonantengraphemen (*K) und Vokalgraphemen (*V) sowie überflüssige Konsonantengrapheme (K+). Bei den Reduktionsvokalen dominieren fehlende (RV-) und fehlerhafte (*RV) Wiedergabe von Reduktionsvokalgraphemen. Andere Abweichungen wie ausgelassene Vokalgrapheme (V-), fehlerhaftes (*VR), fehlendes (VR-) oder überflüssiges vokalisiertes <r> (VR+), überflüssige (V+) und fehlende Vokalgrapheme (V-) sowie fehlerhafte Graphemabfolgen (GV) zeigen eine relativ geringe Fehlerhäufigkeit.

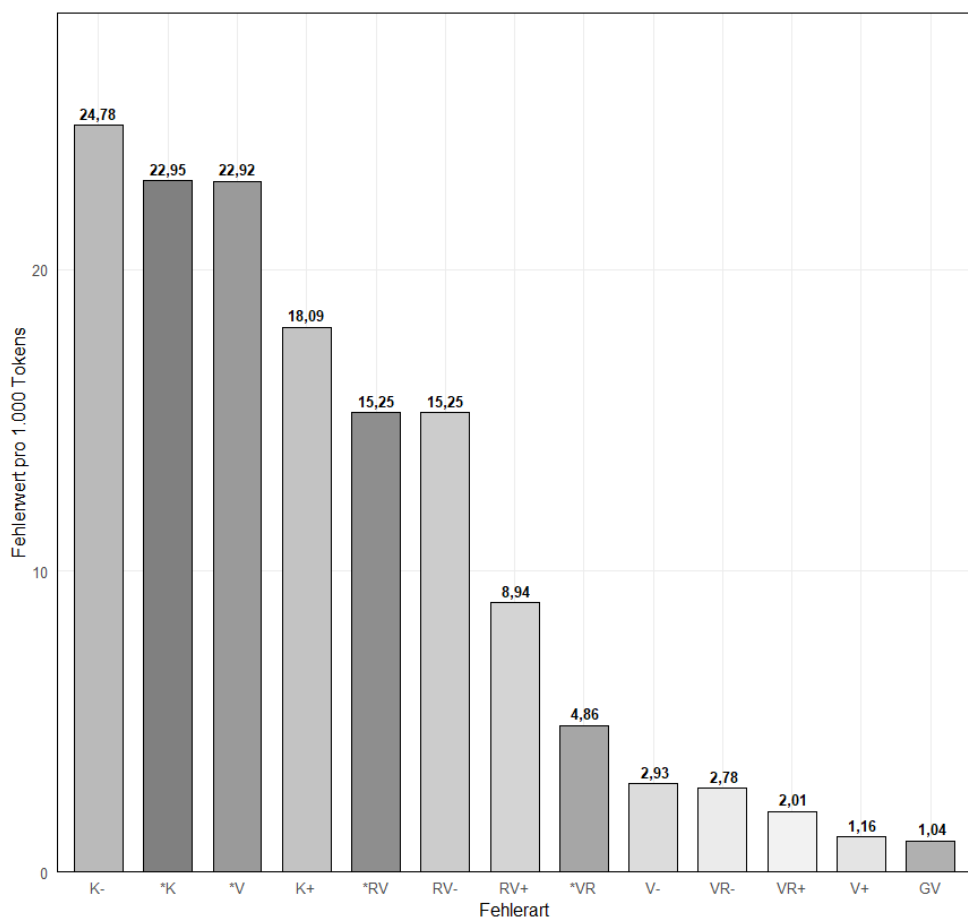


Abbildung 16: PGK-Fehlerverteilung nach Fehlerart; Fehlerwerte pro 1.000 Tokens

Tabelle 20 schlüsselt die Abweichungen nach den vier grundlegenden Fehlertypen auf: der fehlerhaften Auswahl, dem Weglassen, dem Hinzufügen und der falschen Reihenfolge von Graphemen. Dabei werden wegen Platzeinschränkung für jede Kategorie die repräsentativsten Abweichungstypen in der Form [Graph_Fehler] – [Graph_ZH] dargestellt.

Tabelle 20: Typische PGK-Abweichungen nach Unterkategorien; Fehlerwert pro 1.000 Tokens

Abweichungs- klassen	Kategorien	Graph_Fehler – Graph_ZH	Anzahl	Fehlerwerte pro 1.000 Tokens
fehlerhafte Auswahl von Graphemen	falsche Vokalgrapheme (*V)	<ie> – <e>	81	2,56
		<ü> – <e>	47	1,48
		<ei> – <e>	45	1,42
		<i> – <e>	31	0,98
		<e> – <ie>	27	0,85
	falsche Konsonan- tengrapheme (*K)	<n> – <m>	77	2,43
		<m> – <n>	68	2,15
		<t> – <d>	29	0,92
		<n> – <l>	28	0,88
		<l> – <n>	26	0,82
	falscher Reduktionsvokal (*RV)	<e> – <er>	333	10,51
		<er> – <e>	105	3,32
		<u> – <e>	9	0,28
		<a> – <e>	7	0,22
		<r> – <er>	6	0,19
	falsches vokalisiertes <r> (*VR)	<l> – <r>	92	2,90
		<n> – <r>	38	1,20
		<ll> – <r>	9	0,28
		<er> – <r>	5	0,16
		<ch> – <r>	4	0,13
Summe der fehlerhaften Auswahl von Graphemen			2090	65,98
Fehlende Grapheme	fehlende Konsonan- tengrapheme K-	<ø> – <n>	371	11,71
		<ø> – <t>	87	2,75
		<ø> – <l>	68	2,15
		<ø> – <s>	56	1,77
		<ø> – <r>	44	1,39
	fehlende Reduktions- vokalgrapheme (RV-)	<ø> – <e>	386	12,19
		<ø> – <er>	90	2,84
		<ø> – <r>	2	0,06

(Fortsetzung Tabelle 20)

Abweichungs- klassen	Kategorien	Graph_Fehler – Graph_ZH	Anzahl	Fehlerwerte pro 1.000 Tokens
	fehlende Vokalgrapheme (V-)	<ø> – <a>	18	0,57
		<ø> – <u>	13	0,41
		<ø> – <i>	12	0,38
		<ø> – <e>	10	0,32
		<ø> – <ie>	10	0,32
	fehlendes vokalisiertes <r> (VR-)	<ø> – <r>	88	2,78
Summe der fehlenden Grapheme			1449	45,75
„überflüssige“ Grapheme	überflüssige Konso- nantengrapheme (K+)	<n> – <ø>	383	12,09
		<t> – <ø>	51	1,61
		<l> – <ø>	38	1,20
		<d> – <ø>	25	0,79
		<s> – <ø>	16	0,51
	überflüssige Reduk- tionsvokalgrapheme (RV+)	<e> – <ø>	254	8,02
		<er> – <ø>	27	0,85
	überflüssige Vokalgrapheme (V+)	<a> – <ø>	10	0,32
		<i> – <ø>	9	0,28
		<ei> – <ø>	4	0,13
		<ie> – <ø>	3	0,09
		<e> – <ø>	2	0,06
	überflüssiges vokalisiertes <r> (VR+)	<r> – <ø>	63	1,99
Summe der überflüssigen Grapheme			957	30,21
fehlerhafte Abfolge	Graphem- vertauschung (GV)	<ra> – <ar>	10	0,32
		<re> – <er>	6	0,19
Summe der Graphemvertauschungen			33	1,04

Die Fehlerverteilungen in der Tabelle 20 beweisen, dass die Lernenden vor allem Schwierigkeiten mit der korrekten Zuordnung von phonetischen Lauten zu den graphematischen Entsprechungen haben. Dabei haben die Lernenden besonders Schwie-

rigkeiten bei der Unterscheidung der Reduktionsvokale /ə/ (<e>) und /ɐ/ (<er>) (n = 1.249; 39,43 Fehler pro 1.000 Tokens). Außerdem wird das Graphem <l> von den Lernenden häufig als graphematische Entsprechung eines vokalisierten <r> verschriftlicht. Die fehlerhafte Auswahl von Graphemen betrifft sowohl Konsonanten als auch Vokale. Bei den Vokalgraphemen bereitet die Verschriftlichung der *e*-Laute den Lernenden besondere Schwierigkeiten. Zudem stellt die Verwechslung der Konsonanten <m> und <n> eine wichtige Fehlerquelle dar. Bei fehlenden Graphemen dominieren die Auslassung des Reduktionsvokalgraphems <e> und des Konsonantengraphems <n>. Bei überflüssigen Graphemen sind es ebenfalls <n> und <e>, die am häufigsten fälschlicherweise hinzugefügt werden.

Die hohe Fehlerdichte im PGK-Bereich verdeutlicht die grundlegenden Herausforderungen, die das deutsche Laut-Buchstaben-System für chinesische Deutschlernende birgt. Die Schwierigkeiten mit deutschen Reduktionsvokalen (<e>, <er>) und dem vokalisierten <r> sind typisch für Lernende, deren L1 (Mandarin-Chinesisch) keine phonetisch vergleichbar reduzierten Laute kennt. Deren Perzeption und korrekte Verschriftung ist oft von der Sprechgeschwindigkeit und dem Kontext abhängig. In diesem Sinne sind die deutschen Laute anspruchsvoll für chinesische Deutschlernende. Die häufige Substitution des vokalisierten <r> durch <l> sowie <m> und <n> mag auf eine auditive Ähnlichkeitswahrnehmung in einer Diktataufgabe zurückzuführen sein. Die Daten aus der Annotationsebene [Son_Gra-Ab] deuten jedoch darauf hin, dass die Verwechslung der Nasale <m> und <n> zu einem erheblichen Teil (79,4 %) nicht rein phonologisch ist, sondern grammatisch (Deklinationsfehler) bedingt sein kann.

Das häufige Auslassen oder Hinzufügen von <n> und <e> ist nicht zufällig, sondern lässt sich direkt durch die phonotaktischen Strukturen des Chinesischen begründen. Die hohe Frequenz von silbenauslautendem /n/ im Chinesischen, wo es neben /ŋ/ der einzig mögliche konsonantische Auslaut ist (vgl. Kapitel 2.3.2), führt bei den Lernenden zu einer Übergeneralisierung und somit zur Hinzufügung von <n> (T. Liu, 2015, 81 f.; M. Wang, 2010, 43 f.). Die häufige Epenthese (überflüssiges <e>) ist eine L1-basierte Strategie zur Auflösung von im Chinesischen unüblichen Konsonantenclustern, die sich von der Aussprache auf die Schrift überträgt (vgl. Hunold, 2009, S. 154, 2021, S. 8; T. Liu, 2015, S. 82; Tang, 2003, S. 63). Die Verwechslung von <l> und <n> könnte bei einigen Lernenden auf dialektale Einflüsse ihrer chinesischen Herkunftsregion (vor allem Südchina) zurückzuführen sein, in denen /n/ und /l/ nicht immer klar unterschieden werden (vgl. F. Li, 2014, S. 16; M. Wang, 1988, S. 80).

Obwohl Graphemvertauschungen (GV) im gesamten Korpus nur selten auftreten, verdeutlicht die qualitative Analyse ein sehr spezifisches Muster: Sämtliche beobachteten Fälle betreffen die Abfolge von Vokal und <r>, insbesondere in Kontexten, in denen standardmäßig ein vokalisiertes <r> zu erwarten wäre (z. B. <ar>, <er>). Die Analyse der Diktataufnahmen liefert hierfür eine plausible Erklärung: Der Vorleser realisierte das <r> in diesen Positionen nicht als Vokal (/ɐ/), sondern als deutlich hörbares, „gerolltes“ *r* mit mehreren Vibrationen der Zunge. Diese für die Lernenden vermutlich unerwartete und sehr saliente konsonantische Aussprache führte dazu, dass sie den gehörten Konsonanten vor dem Vokal verschriftlichten, was die Graphemver-

tauschung (z. B. *<Graten> statt <Garten>) zur Folge hatte. Solche Fehlschreibungen in dieser Kategorie sind vielmehr eine direkte Reaktion auf eine spezifische, nicht standardisierte Variante im auditiven Input.

Interessanterweise weichen die hier präsentierten Ergebnisse zur Verteilung und Art der PGK-Fehler teilweise von den Befunden Tangs (2003) ab, der ebenfalls schriftsprachliche Abweichungen chinesischer Deutschlernender untersuchte, jedoch auf Basis von Aufsätzen. In Tangs Studie dominierten unter den „phonographematischen Schreibabweichungen“ – einer Kategorie, die im Wesentlichen der PGK-Ebene im DeChiLKo entspricht – die „fehlenden Grapheme“, gefolgt von „fehlerhafter Wahl von Graphemen“, während „überflüssige Grapheme“ und „fehlerhafte Graphemabfolge“ seltener auftraten (vgl. Tang, 2003, S. 63), was der Rangfolge im DeChiLKo zwar grob ähnelt, jedoch in den spezifischen Ausprägungen Unterschiede zeigt. So identifizierte Tang (ebd.) bei der fehlerhaften Wahl von Graphemen die Substitutionen <l> für <n>, <e> für <ie> sowie <u> für <ü> als besonders auffällig. Im DeChiLKo hingegen sind es vor allem die Verwechslungen zwischen den Reduktionsvokalgraphemen <e> und <er> sowie die Fehlschreibung <ie> für <e>, die im Vordergrund stehen. Auch bei den fehlenden Graphemen lassen sich Divergenzen beobachten: Während Tang (ebd.) häufige Auslassungen von <t>, <n> und <l> feststellte, tendieren die Lernenden im DeChiLKo primär dazu, das Reduktionsvokalgraphem <e>, das vokalisierte <r> sowie das Konsonantengraphem <n> wegzulassen. Im Bereich der überflüssigen Grapheme beobachtete Tang (ebd.) die häufigste Hinzufügung des Konsonantengraphems <s>, wohingegen im DeChiLKo das Konsonantengraphem <n> am häufigsten überflüssigerweise auftritt. Diese Unterschiede mögen auf die bereits erwähnte Aufgabenabhängigkeit orthographischer Fehlerprofile hinauslaufen: Die spezifischen Anforderungen von Diktaten im DeChiLKo im Vergleich zu freien Aufsatzproduktionen in Tangs Studie könnten die unterschiedlichen Fehlerschwerpunkte bedingen. Darüber hinaus weist Tang (vgl. 2003, S. 172) darauf hin, dass 16,3 % der Schreibabweichungen durch fehlerhafte Aussprache bedingt sein können, die ihrerseits muttersprachlich beeinflusst ist. Außerdem betont Tang, dass Abweichungen durch die lautliche Struktur des Zielphonems selbst oder durch Unkenntnis spezifischer orthographischer Regeln erklärt werden (vgl. Tang, 2003, 122 f.).

Zusammenfassend ist festzuhalten, dass die grundlegenden Phonem-Graphem-Korrespondenzen in der L3 für chinesische Deutschlernende eine erhebliche Herausforderung darstellen. Diese sind sowohl durch strukturelle Unterschiede im sprachlichen Repertoire der Lernenden (insbesondere L1 Chinesisch/*Pinyin* im Vergleich zur L3 Deutsch), als auch durch die inhärente Komplexität des deutschen Laut-Buchstaben-Systems bedingt.

4.2.3 Fehler bei der Markierung von Vokalquantität [Graph_VQM]

Die korrekte orthographische Wiedergabe der Vokalquantität, d. h. der Unterscheidung zwischen langen/gespannten und kurzen/ungespannten Vokalen, stellt ein zentrales und oft fehlerträchtiges Element der deutschen Orthographie dar, das in zahlreichen Studien zum Rechtschreiberwerb als prominenter Untersuchungsgegenstand

gilt (vgl. u. a. Fix, 2002; H.-G. Müller & Schroeder, 2022; Ruppert & Hanulíková, 2022). In etablierten Diagnoseverfahren wird die Relevanz der Vokalquantitätsmarkierung für die Erfassung orthographischer Kompetenzen durch die spezifischen Fehlerkategorien oder Lupenstellen berücksichtigt. So wird beispielsweise in der **AFRA** die Kennzeichnung der Vokalquantität als Sonderbereich der phonologischen Ebene behandelt. Unter dieser Kategorie werden die Fehlschreibungen für das lang gesprochene und betonte /i:/ (LI), für lange Vokale (LV) sowie für Kurzvokale (KV) erfasst, wobei jeweils nach der Häufigkeit des Vorkommens des Zielgraphems nach Minderheit und Mehrheit unterschieden wird (vgl. Herné & Naumann, 2016, S. 10–12). Die HSP betrachtet die Anwendung von Längenzeichen und Kürzezeichen als wichtige „Lupenstellen“ innerhalb der Rechtschreibstrategie (May, 2013, S. 15). Auch die OLFA widmet der Vokalquantitätsmarkierung breiten Raum und kategorisiert detailliert verschiedene Fehlschreibungen unter den Fehlerkategorien 07 bis 11. Dazu gehören „Einfachschreibung für Konsonantenverdopplung“; „Konsonantenverdopplung für Einfachschreibung nach Kurzvokal“; „einfache Vokalschreibung für markierte Länge und <i> für <ie> bei langem /i:/“; „markierte Längen- für Einfachschreibung bei Langvokal; Konsonantenverdopplung nach Langvokal“; „Konsonant oder am Morphemanfang“; „markierte Vokallänge bei Kurzvokal, auch <ie> für kurzes <i>“ (Thomé & Thomé, 2020, S. 20–22). Diese differenzierte Erfassung in den Testverfahren drückt die Komplexität und Wichtigkeit dieses Bereichs für die Entwicklung orthographischer Kompetenz aus.

In der vorliegenden Studie werden Abweichungen bei der Markierung von Vokalquantität auf der Annotationsebene [Graph_VQM] erfasst. Dabei werden Verstöße gegen die bezeichnete oder unbezeichnete Länge und Kürze der Vokale als VM+ (markierte Längenschreibung für den unbezeichneten Langvokal), VM- (Einfachvokalschreibung für den bezeichneten Langvokal), *VM (markierte Längenschreibung für den Kurzvokal), VM° (falsche Bezeichnung für den Langvokal), KD+ (Konsonantenverdopplung für den Kurzvokal), KD- (Einfachkonsonanz für Konsonantenverdopplung) und *KD (Konsonantenverdopplung für den Langvokal, nach dem Konsonanten oder am Morphemanfang) annotiert. Als markierte Längenschreibung werden Verdoppelung von Vokalgraphemen, Einsetzung des Dehnungs-<h> sowie die Graphemfolge <ie> berücksichtigt. Unter doppelten Konsonantengraphemen werden einfache Verdopplungen wie <bb>, <dd>, <ff>, <gg>, <ll>, <mm>, <nn>, <pp>, <rr>, <ss>, <tt> sowie die beiden Sonderfälle <ck> und <tz> berücksichtigt, die in der deutschen Orthographie positionsbedingt anstelle von <k> und <z> nach Kurzvokalen stehen.

Im DeChiLKo-Korpus wurden insgesamt 439 Abweichungen bei der Kürzen- und Längenmarkierung in 207 Lernertexten identifiziert, was einem Fehlerwert von 13,86 pro 1.000 Tokens entspricht.

Tabelle 21: Fehlerwert der VQM-Abweichungen pro 1.000 Tokens

VQM-Abweichungen insgesamt	Summe aller Tokens	VQM-Annotationen pro 1.000 Tokens
439	31.674	13,86

Eine detaillierte Analyse der Fehlerverteilung nach Fehlerarten (vgl. Abbildung 17) präzisiert, dass die Lernenden am häufigsten bezeichnete Langvokale als Einfachvokale verschriftlichen (VM-: 4,58/1.000 Tokens). An zweiter Stelle steht die Einfachschreibung von Konsonanten bei erforderlicher Konsonantenverdopplung (KD-: 3,79/1.000 Tokens), gefolgt von der fehlerhaften Konsonantenverdopplung nach Langvokal, Konsonant oder am Morphemanfang (*KD: 2,46/1.000 Tokens). Auch die markierte Längenschreibung bei ungespannten Vokalen (*VM: 1,70/1.000 Tokens) stellt eine nennenswerte Fehlerquelle dar. Andere Fehlerarten wie die Verdopplung von Konsonanten für den unmarkierten Kurzvokal (KD+: 0,88/1.000 Tokens), die überflüssige Markierung eines unbezeichneten Langvokals (VM+: 0,41/1.000 Tokens) oder eine falsche Bezeichnung für den Langvokal (VM*: 0,03/1.000 Tokens) sind hingegen selten.

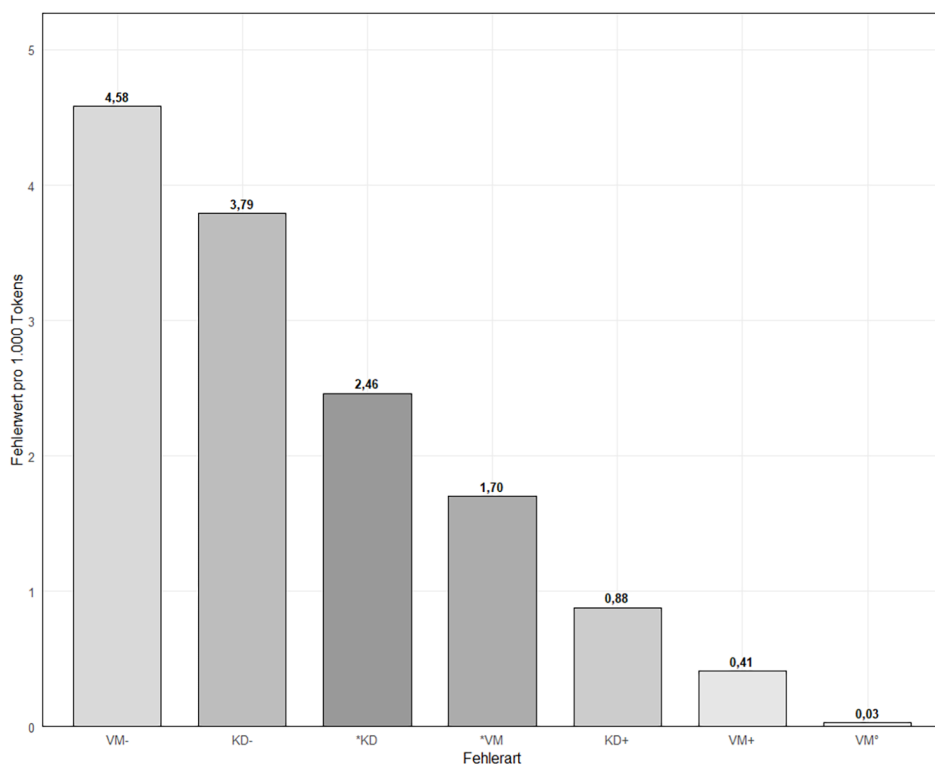


Abbildung 17: VQM-Fehlerverteilung nach Fehlerart; Fehlerwerte pro 1.000 Tokens

Zusammenfassend lässt sich eine deutliche Tendenz zur Untermarkierung der Vokalquantität feststellen. Diese manifestiert sich primär in der Auslassung von Vokallängenmarkierung (VM-) und Konsonantenverdopplung (KD-). Als ursächlich hierfür kann eine unzureichende perzeptive Differenzierung der oft nur subtilen phonetischen Unterschiede zwischen gespannten und ungespannten Vokalen im Deutschen angenommen werden. Diese mangelnde phonologische Bewusstheit kann konsequent zu einer

fehlerhaften Anwendung der entsprechenden orthographischen Markierungsregeln führen (vgl. Hunold, 2021, S. 8). Der geringe Fehlerwert bei überflüssiger oder gänzlich falscher Markierung deutet jedoch darauf hin, dass ein grundlegendes Bewusstsein für die Existenz von Längen- und Kürzenmarkierungen im Deutschen vorhanden ist, die Anwendung aber unsicher bleibt.

Nachdem die allgemeine Fehlerverteilung und -häufigkeit in der Kategorie Vokalquantitätsmarkierung erläutert wurde, wird nun eine detaillierte Untersuchung der spezifischen Fehlergrapheme durchgeführt. Hierfür werden die Annotationsebenen [Graph_Fehler] und [Graph_ZH] einbezogen. Im gesamten DeChiLKo sind insgesamt 54 unterschiedliche Abweichungen zu beobachten, bei 24 davon handelt es sich um Dehnungsgraphien, und 30 Abweichungen kommen bei Schärfungsgraphien vor. In der folgenden Tabelle 22 werden für jede Fehlerart die 10 repräsentativsten Abweichungstypen dargestellt.

Tabelle 22: Typische Abweichungen in der VQM-Kategorie

Nr.	Fehlerart	Graph_Fehler-Graph_ZH	Summe	Fehlerwert (pro 1.000 Tokens)
1	VM-	<i> - <ie>	47	1,41
2	*KD	<mm> - <m>	36	1,12
3	VM-	<e> - <eh>	35	1,07
4	VM-	<o> - <oh>	30	0,91
5	*VM	<eh> - <e>	28	0,84
6	KD-	<n> - <nn>	25	0,75
7	KD-	<s> - <ss>	25	0,75
8	KD-	<l> - <ll>	22	0,67
9	VM-	<a> - <ah>	21	0,64
10	VM-	<ie> - <i>	19	0,57

Ähnlich wie für andere Lernergruppen mit L1 wie Russisch, Türkisch usw. ist die Vokallänge für chinesische Deutschlernende auch kein distinktives phonologisches Merkmal in ihrer L1. Die Daten in Tabelle 22 stützen die Beobachtungen von Ruppert und Hanulíková (2022). Die Autorinnen konstatieren, dass eine primäre Sozialisation in Sprachen, die nicht systematisch zwischen gespannten und ungespannten Vokalen differenzieren, zu erheblichen Diskriminierungsschwierigkeiten im Deutschen führt. Dies schlägt sich konsequent in entsprechenden Problemen bei der orthographischen Umsetzung nieder (Ruppert & Hanulíková, 2022, S. 70).

Im Gegensatz dazu scheinen die Lernenden weniger Probleme mit der Vokalverdopplung als Längenkennzeichen zu haben. Dies könnte auf einen positiven Transfer-effekt aus der ersten Fremdsprache Englisch zurückzuführen sein, wo Vokaldopplun-

gen (z. B. *see, too*) ebenfalls zur Markierung von Vokallänge dienen. Obwohl die genauen Regeln und Funktionen nicht identisch sind, könnte die Vertrautheit mit dem Prinzip der Vokaldopplung aus dem Englischen den Erwerb im Deutschen erleichtern und die relativ geringe Fehlerrate in diesem spezifischen Bereich erklären.

Als Dehnungsgraphie unterliegt das Dehnungs-*<h>* einigen spezifischen phonotaktischen Bedingungen. Anders als die Vokalverdopplung findet es weder in der L1 Chinesisch noch in der ersten Fremdsprache Englisch eine direkte Entsprechung. Die Lernenden können hier kaum auf vorhandenes Wissen aus ihrem multikompetenten Repertoire zurückgreifen.

Besonders auffällig ist die hohe Fehlerfrequenz bei der Graphemauswahl für den Langvokal /i:/. Mit einem Fehlerwert von 1,35 pro 1.000 Tokens stellt die Substitution von *<ie>* durch *<i>* die häufigste Abweichung in diesem Bereich dar. Dieser Befund kontrastiert deutlich mit der allgemeinen Sprachstatistik des Deutschen: Tabelle 23 zeigt, dass *<ie>* mit einer Häufigkeit von 35,88 Treffern pro 1 Mio. Tokens die mit Abstand dominanteste Markierung für die Vokalquantität im Referenzkorpus (DWDS Kernkorpus 21) darstellt.¹⁸

Dass die Lernenden trotz dieser hohen Frequenz in der geschriebenen Standardsprache gerade hier die meisten Fehler machen, unterstreicht, dass statistische Häufigkeit nicht zwangsläufig zu orthographischer Korrektheit führt. Zwar bildet ein allgemeines Referenzkorpus wie das DWDS-Kernkorpus 21 den spezifischen, didaktisch aufbereiteten Input chinesischer Deutschlernender nicht eins zu eins ab, es dient hier jedoch als exemplarischer Indikator für die grundlegende statistische Verteilung im deutschen Schriftsystem. Es ist einzuräumen, dass die hier genutzte Token-Frequenz durch hochfrequente Funktionswörter (z. B. *sie, die*) beeinflusst wird. Dennoch bildet sie die tatsächliche Input-Exposition der Lernenden ab, deren Erwartungshorizonte durch die visuelle Präsenz dieser Grapheme im Alltag geprägt werden. Die Beobachtung, dass Frequenz allein nicht ausreicht, deckt sich zudem mit Erkenntnissen aus dem L1-Schrifterwerb: Auch deutschsprachige Grundschulkinder zeigen bei der Markierung des Typs *<ie>* große Schwierigkeiten, da diese Schreibung eine explizite Regelkenntnis erfordert, die über die reine Laut-Buchstaben-Zuordnung hinausgeht (vgl. Fay, 2010).

¹⁸ Für die Abfrage im DWDS-Kernkorpus 21 wurde versucht, die Kennzeichnungen der langen Vokale im Referenzkorpus möglichst präzise zu identifizieren, während andere, morphologisch oder phonetisch abweichende Graphemfolgen von den Ergebnissen ausgeschlossen werden sollten. Für die Ermittlung der Frequenz des Graphems *<ie>* wurde eine systeminterne Zeichenketten-Suche durchgeführt. Hierbei wurden morphosyntaktisch bedingte Sonderfälle wie unbetonte *<ie>*-Verbindungen in bestimmten Suffixen (z. B. *Familie, Studie, Aktie*) manuell aus den Trefferlisten bereinigt. Es ist jedoch anzumerken, dass vereinzelte heterosyllabische Fremdwortstrukturen (wie etwa bei „Klient“), bei denen *<ie>* keinen Langvokal /i:/ repräsentiert, aufgrund der algorithmischen Limitationen nicht vollständig ausgeschlossen werden konnten und als marginale Unschärfe in den Gesamtwerten enthalten bleiben. Analog dazu wurden bei der Suche nach der Graphemverdopplung *<ee>* gezielt Treffer wie „beeinflussen“ ausgeschlossen. Ebenso wurden Wortformen wie „sehen“ oder „Nahe“ von den Treffern entfernt, da in diesen Fällen das Graphem *<h>* silbeninitial ist und somit nicht als Markierung langer Vokale betrachtet wird.

Tabelle 23: Treffer der unterschiedlichen Kennzeichnungen der langen Vokale im DWDS-Kernkorpus 21 (Quelle:DWDS Kernkorpus 21)

Nr.	lange Vokale	Kennzeichnungen der langen Vokale	Treffer im Korpus (Token)	Häufigkeit pro 1 Mio. Tokens
1	/i:/	<ie>	552.722	35,88
2	/e:/	<eh>	132.941	8,63
3	/a:/	<ah>	93.359	6,06
4	/i:/	<ih>	69.483	4,51
5	/ɔ:/	<oh>	41.865	2,72
6	/ɛ:/	<äh>	36.710	2,38
7	/y:/	<üh>	35.652	2,31
8	/e:/	<ee>	22.529	1,46
9	/ɔ:/	<oo>	20.444	1,33
10	/a:/	<aa>	17.060	1,11
11	/i:/	<ieh>	12.482	0,81
12	/u:/	<uh>	9.047	0,59
13	/y:/	<öhh>	6.531	0,42

Die Abweichungen bei der Konsonantenverdopplung zeigen zudem, dass vor allem nasale und alveolare Konsonanten wie /n/, /m/ und /l/ betroffen sind. Fehler bei der Markierung von Schärfungsgraphemen, wie das Auslassen von <ss>, kommen ebenfalls häufig vor.

In dieser Annotationsebene [Graph_VQM] werden die Abweichungen primär unter dem Aspekt der direkten Graphem-Phonem-Zuordnung im Hinblick auf Vokalquantität betrachtet. Morphologisch bedingte Schreibweisen, die ebenfalls die Vokalquantität betreffen (z. B. Stammkonstanz bei „Schwimmbad“ durch morphologische Gelenkschreibung oder bei „Fahrt“ durch morphologische Dehnungsschreibung), werden in den entsprechenden morphologischen Annotationsebenen (siehe Kapitel 4.4) detailliert untersucht.

4.2.4 Fehler bei der s-Schreibung [Graph_s-Schreibung]

Die Annotationsebene [Graph_s-Schreibung] thematisiert die Systematik der Schreibung der beiden s-Laute [z] (stimmhaftes s) und [s] (stimmloses s) mit den Graphemen <s> und <ß> sowie der Graphemfolge <ss>. Eisenberg (vgl. 1995: 180) bezeichnet die s-Schreibung als einen schwierigen Bereich der deutschen Orthographie, da die historisch gewachsene, nicht immer eindeutige Zuordnung der Laute zu den drei Graphemen eine potenzielle Fehlerquelle für Lernende darstellt. Die Rechtschreibreform von 1996 zielte darauf ab, hier mehr Systematik und Konsistenz zu schaffen: Der

Wechsel von <ss> zu <ß> nach kurzem Vokal wurde aufgehoben und <ss> konsequent für die Schreibung nach Kurzvokal etabliert (z. B. *er muss* statt *er muß*). Das Graphem <ß> dient seither der Kennzeichnung eines vorausgehenden langen Vokals oder Diphthongs vor dem stimmlosen s-Laut (z. B. *Maß*, *draußen*). Zudem wurde die einheitliche Schreibung der unterordnenden Konjunktion *dass* anstelle von *daß* eingeführt (vgl. Heller, 1996).

Obwohl die Rechtschreibreform nunmehr drei Jahrzehnte zurückliegt, ist die normgerechte s-Schreibung im aktuellen Sprachgebrauch nicht immer vollständig etabliert. So finden sich beispielsweise im DWDS Kernkorpus 21 immer noch Belege für die veraltete Schreibweise *muß* in Zeitungsartikeln (wie z. B. *Die Zeit*) aus dem Zeitraum 2000 bis 2010. Angesichts des nach der Rechtschreibreform von 1996 bestehenden Übergangszeitraums ist jedoch zu berücksichtigen, dass ein Teil dieser Schreibungen im jeweiligen Entstehungszeitraum möglicherweise noch als normgerecht galt. Gleichwohl untergräbt die Persistenz solcher Formen in als normgebend wahrgenommenen Textsorten die Verbindlichkeit der offiziellen Orthographieregeln für den Spracherwerb. Aus der Perspektive der CDST wirkt ein solcher inkonsistenter Input für den Spracherwerb als ein externer Faktor, der die Ausbildung eines stabilen Attraktorzustandes für die korrekte orthographische Form verhindern kann (vgl. N. C. Ellis & Larsen-Freeman, 2006; Larsen-Freeman, 2017). Hierbei ist einzuräumen, dass für Lernende außerhalb des deutschsprachigen Raums der primäre Input meist durch hochgradig normierte Lehrwerke gesteuert wird. Dennoch könnte die Persistenz veralteter Formen in prestigeträchtigen Quellen bei fortgeschrittenen Lernenden, die mit authentischen Texten arbeiten, als konkurrierender Attraktor wirken und die Phase der Unsicherheit verlängern.

Die Fehleranfälligkeit der s-Schreibung spiegelt sich auch in etablierten Rechtschreibtestverfahren wider. Im OLFA-Test (Thomé & Thomé, 2020: 22) werden Abweichungen in der s-Schreibung (<s> für <ß>, <ß> für <s>, <ss> für <ß> sowie <ß> für <ss>) den Kategorien „Gruppe II (alphabetisch)“ und „Gruppe III (orthographisch)“ zugeordnet. Im AFRA-System (Herné & Naumann, 2016) werden Fehler bei Graphemen <s> und <ß> in der Kategorie „Spezielle Grapheme“ klassifiziert, wobei die Fehler bei <s> als „Spezielle Grapheme (Mehrheit)“ und Fehler bei <ß> als „Spezielle Grapheme (Minderheit)“ bezeichnet werden.

Um die orthographische Leistung chinesischer Deutschlernender im Bereich der s-Schreibung präzise zu erfassen, wurde in der vorliegenden Arbeit auf der Annotationsebene [Graph_S-Schreibung] besondere Aufmerksamkeit auf die Auswahl der Grapheme <s>, <ss> und <ß> gelegt. Abweichungen wurden dabei in drei Kategorien markiert: „SßS“ (Verwechslung der Grapheme <s> und <ß>), „SSßS“ (Verwechslung der Grapheme <ss> und <ß>) sowie „ShS“ (Fehler bei der stimmhaften S-Schreibung).

Insgesamt wurden in dieser Kategorie 45 Abweichungen in 40 Lernertexten im gesamten DeChiLko festgestellt, was einem relativ niedrigen Fehlerwert von 1,42 pro 1.000 Tokens entspricht (vgl. Tabelle 24).

Tabelle 24: Fehlerwert der s-Schreibung-Abweichung pro 1.000 Tokens

s-Schreibung-Abweichungen insgesamt	Summe aller Tokens	s-Schreibung-Annotationen pro 1.000 Tokens
45	31.674	1,42

Die Verteilung der Abweichungen nach den drei Hauptfehlerarten (vgl. Abbildung 18) zeigt, dass die Verwechslungen der Grapheme <ss> und <ß> (SSßS) bei der Schreibung des stimmlosen s-Lauts mit einem Fehlerwert von 0,66/1.000 Tokens am häufigsten auftreten. Darauf folgen Abweichungen bei der Schreibung des stimmhaften [z]-Lauts (ShS) mit einem Fehlerwert von 0,47/1.000 Tokens. Verwechslungen zwischen <s> und <ß> (SßS) bei der Verschriftlichung des stimmlosen [s]-Lauts sind hingegen selten.

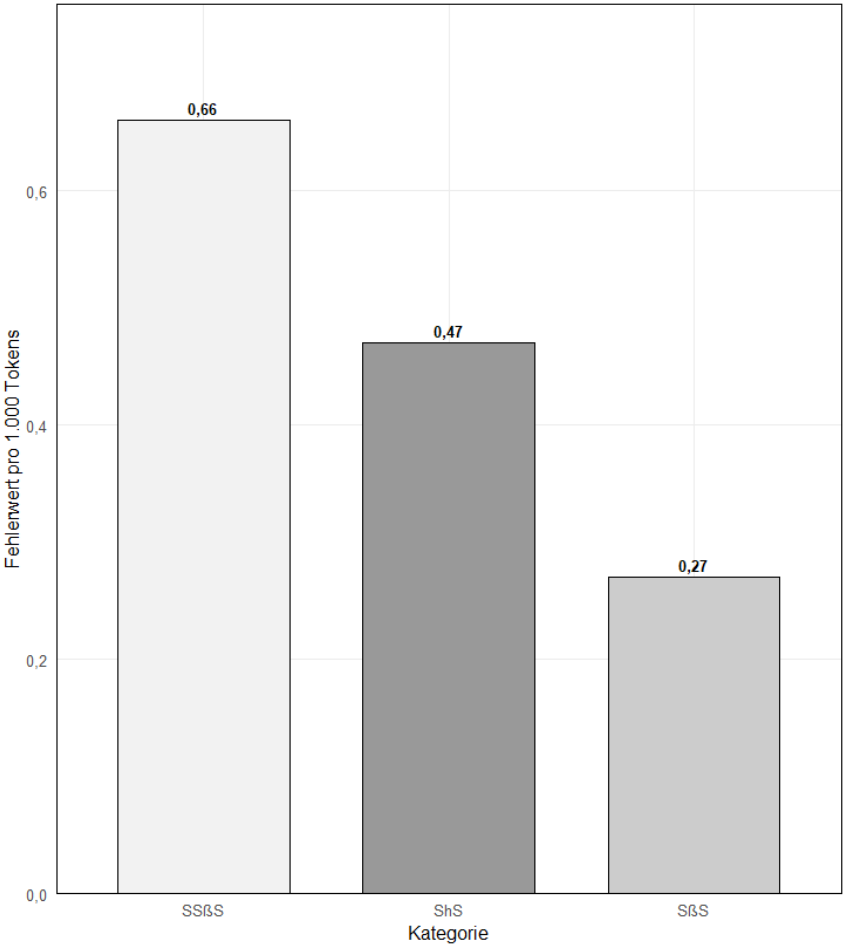


Abbildung 18: Fehlerverteilung der s-Schreibung nach Fehlerarten

Eine genaue Analyse der fünf spezifischen *s*-Schreibung-Fehlertypen (vgl. Abbildung 19) verdeutlicht, dass die Fehlerverteilung <z> für <s> (bzw. <Z> für <S>) bei der Realisierung des stimmhaften [z]-Lauts der häufigste Fehlertyp ist. Bei der Schreibung des stimmlosen [s]-Lauts besteht die häufigste Fehlerart in der Verwechslung zwischen den Graphemen <ss> und <ß>, wobei die Schreibung <ß> für <ss> (0,47/1.000 Tokens) ebenfalls den ersten Platz einnimmt und damit mehr als doppelt so häufig vorkommt wie die umgekehrte Fehlschreibung <ss> für <ß> (0,19/1.000 Tokens). Die Fehlschreibungen <s> für <ß> sowie <ß> für <s> bilden mit Fehlerwerten von 0,19 bzw. 0,09 das Schlussfeld.

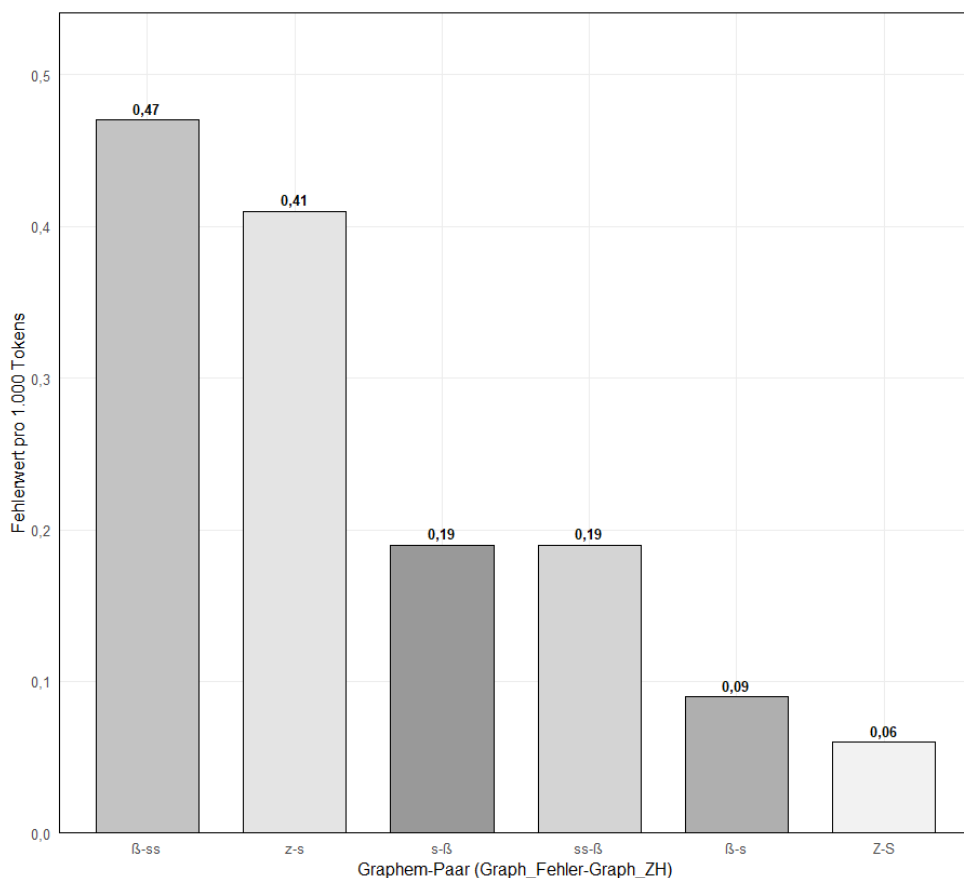


Abbildung 19: Fehlerverteilung im Bereich S-Schreibung nach Fehlertypen ([Graph_Fehler] – [Graph_ZH])

Die Ergebnisse zur *s*-Schreibung im DeChiLKO erörtern ein spezifisches Fehlerprofil für chinesische Deutschlernende, das sich von den Befunden anderer Lernergruppen unterscheidet und im Kontext ihres multikompetenten Repertoires interpretiert werden kann. Der Vergleich mit den Forschungsergebnissen von Thomé (1987) zu orthographischen Fehlern bei Schüler*innen mit Türkisch als Erstsprache verdeutlicht dies. Laut

Thomé machen türkische Schüler*innen die meisten Fehler bei der Verwechslung der Grapheme <s> und <ß> (<s> für <ß>: 5,40 Fehler /1000 Wörter¹⁹; <ß> für <s>: 0,82 Fehler/1000 Wörter). Die Verwechslung zwischen <ss> und <ß> trat in seiner Studie hingegen selten auf (<ß> für <ss>: 0,26; <ss> für <ß>: 0,16/ 1000 Wörter) (vgl. Thomé, 1987: 102).

Diese deutlichen Unterschiede sind zum einen auf den historischen Zeitpunkt der Analyse zurückzuführen: Thomés Untersuchung wurde vor der Rechtschreibreform von 1996 durchgeführt. Da die Unterscheidung von *das* und *daß* (*dass*) eine erhebliche Fehlerquelle darstellte (vgl. Thomé, 1987: 387 f.), stieg die Häufigkeit von Verwechslungen zwischen <s> und <ß> in seiner Untersuchung. Zum anderen spielen die unterschiedlichen L1-Hintergründe eine Rolle. Der im DeChiLKo häufigste Fehlertyp, <z> für <s> bei stimmhaftem /z/, ist bei den von Thomé untersuchten türkischen Lernenden mit einem Fehlerwert von nur 0,52 pro 1000 Wörter deutlich unterrepräsentiert. Thomé (vgl. 1987: 112 f.) führt diese Fehlschreibung bei seinen Probanden einerseits auf den Berliner Dialekt zurück, andererseits auf graphemische Interferenz aus dem Türkischen, wo das stimmhafte *s* /z/ mit dem Graphem <z> verschriftlicht wird. Für die chinesischen Deutschlernenden im DeChiLKo ist diese spezifische L1-Interferenzquelle nicht gegeben. Stattdessen könnte der hohe Fehlerwert bei der Abweichung <z> für <s> darauf hindeuten, dass die Schreibung des stimmhaften /z/ für sie eine besondere Lernhürde darstellt. Im Chinesischen existieren zwar dentale Frikative, jedoch ausschließlich in stimmloser Form (der Konsonant /s/); eine direkte Entsprechung für den stimmhaften *s*-Laut [z] fehlt im chinesischen Konsonantensystem (vgl. Hachenberg, 2003; Yen, 1992). Dies kann zu Schwierigkeiten bei der auditiven Perzeption und Differenzierung von /s/ und /z/ führen. Hachenberg (2003, 110 ff.) stellt in ihrer Studie zu Ausspracheproblemen chinesischer Deutschlernender fest, dass das stimmhafte /z/ häufig durch die dentale Affrikate /tʃ/ (*Pinyin* <z>) ersetzt wird, was eine L1-basierte phonetisch-phonologische Interferenz darstellt, die sich potenziell auch auf die Graphemwahl auswirkt. Die Fehlschreibung <z> für <s> könnte somit einerseits durch den Einfluss der L2 Englisch erklärt werden, da im Englischen das stimmhafte /z/ mit dem Graphem <z> geschrieben wird (z. B. *zoo*, *lazy*). Andererseits könnte auch eine logographemische Ähnlichkeit zwischen dem Laut und dem Buchstaben <z> eine Rolle spielen, insbesondere wenn die korrekte deutsche Graphem-Phonem-Korrespondenz für /z/ (nämlich <s> in vielen Kontexten) noch nicht gefestigt ist.

Die von Thomé (1999, 219 ff.) beschriebenen Erwerbsstufen der *s*-Schreibung bieten einen weiteren Interpretationsrahmen. Die Fehlschreibung <s> für <ß> beim Phonem /s/ gilt als „Grundfehler“ der alphabetischen Phase der inneren Regelbildung (vgl. Thomé, 1999, S. 219). In der nächsten Stufe, der „ß-Erarbeitung“, tritt dann die Übergeneralisierung <ß> für <s> auf (z. B. *<Preiße> statt <Preise>). Diese Fehler können aufgrund ihrer Zufälligkeit als semi-arbiträre Schreibung bezeichnet werden, d. h. als Schreibungen, bei denen kaum Anhaltspunkte für die Verwendung von <ß> erkennbar sind und die daher fast willkürlich wirken (vgl. Thomé, 1999, S. 219). Die Zunahme von

19 Die Originalwerte bei Thomé (1987) sind in Fehlern pro 500 Wörter angegeben. Für eine bessere Vergleichbarkeit mit den übrigen in dieser Studie verwendeten Werten wurden die Angaben auf 1.000 Wörter umgerechnet.

<ß> für <ss> sowie das seltenere Auftreten von <ss> für <ß> deutet auf die Entwicklungsstufe II hin, in der lautliche Bedingungen zwar berücksichtigt, morphematische Überlegungen (Stammkonstanz) aber noch vernachlässigt werden. In Stufe III werden die morphematische Orientierung und die lautliche Kontrolle miteinander in Einklang gebracht. Besonders auffällig in dieser Phase ist das korrekte Schreiben des stimmhaften /z/ in der Pluralbildung, wie etwa bei der korrekten Schreibung von <Preise> anstelle von *<Preiße> (vgl. Thomé, 1999). Die Fehlerverteilung im DeChiLKo, insbesondere die Dominanz von <ß> für <ss> über <ss> für <ß>, könnte auf den ersten Blick darauf hindeuten, dass sich viele chinesische Lernende in einer Art „Erwerbsstufe II“ befinden. Die Daten aus den Subkorpora im DeChiLKo (Abbildung 20) zeigen zwar Tendenzen, die bei oberflächlicher Betrachtung mit Thomés Stufenmodell des s-Schreibung-Erwerbs vereinbar scheinen: Der „Grundfehler“ <ß> für <s> tritt ausschließlich im Erwerbskorpora bei Anfängern auf. Anstatt eine lineare Stufe anzunehmen, kann dieses gehäufte Auftreten eines spezifischen Fehlertyps im Sinne der CDST als ein temporärer Attraktorzustand verstanden werden (vgl. Bulut, 2018, S. 52) und dies bedeutet nicht zwangsläufig, dass sie in späteren Phasen vollständig verschwinden oder dass alle Lernenden eine identische, starre Abfolge durchlaufen. Die im Prüfungskorpora weiterhin zu beobachtenden hohen Fehlerhäufigkeiten bei der Verwechslung zwischen <s> und <ß> sowie <ss> und <ß> deuten darauf hin, dass auch fortgeschrittenere Lernende Schwierigkeiten haben, die korrekte Anwendung des <ß>-Graphems zu meistern. Bemerkenswerterweise zeigt die Schreibung des stimmhaften /z/, die Thomé (1999) einer fortgeschrittenen Stufe (Stufe III) zuordnet, im Erwerbskorpora einen tendenziell niedrigeren Fehlerwert als im Prüfungskorpora. Dieses Muster widerspricht einer einfachen Stufenlogik und stützt die Annahme, dass der Erwerb der s-Schreibung kein abgeschlossener, sequenzieller Prozess ist. Die Lernenden bewegen sich nicht strikt von einer definierten Stufe zur nächsten, ihr orthographisches System unterliegt vielmehr einem kontinuierlichen Wandel, in dem verschiedene Wissensbestände und Strategien dynamisch interagieren. Somit können spezifische Fehlertypen oder die ihnen zugrundeliegenden Verarbeitungsschwierigkeiten prinzipiell in jedem Lernstadium auftreten, auch wenn ihre Frequenz je nach Lernstand und Kontext variieren mag.

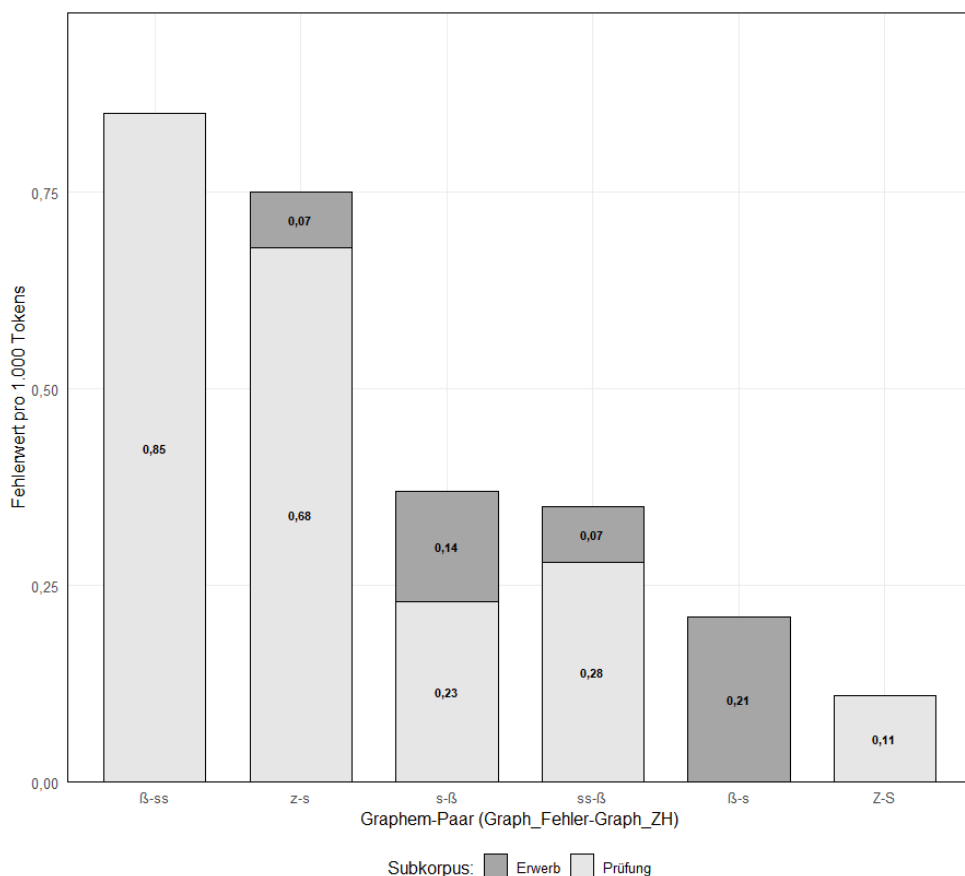


Abbildung 20: Fehlerverteilung der s-Schreibungen nach Subkorpora; Fehlerwerte pro 1.000 Tokens

Zusammenfassend lässt sich festhalten, dass die s-Schreibung für chinesische Deutschlernende quantitativ eine geringere Fehlerquelle darstellt als andere phonographische Bereiche und dass die qualitativen Fehler auf spezifische, nicht lineare Lernprozesse und Interaktionen im multi-kompetenten System hinweisen. Insbesondere die Herausforderung, das im Deutschen bedeutungsunterscheidende Merkmal der Stimmhaftigkeit bei s-Lauten orthographisch korrekt umzusetzen, scheint vor dem Hintergrund der L1 Chinesisch und möglicher L2-Einflüsse eine zentrale Rolle zu spielen.

4.2.5 Fehler bei der Schreibung von „speziellen Graphemen“ [Graph_SG]

Der Begriff „spezielle Grapheme“ wird in Anlehnung an die AFRA (Aachener förderdiagnostische Rechtschreibfehler-Analyse) verwendet (Herné & Naumann, 2016: 9 f.)²⁰.

Spezielle Grapheme (SG) sind fehleranfällige Schriftzeichen, die genau ein Phonem repräsentieren. Sie bereiten aus mehreren Gründen Schwierigkeiten:

- a) Einige von ihnen bestehen aus zwei bzw. drei Buchstaben (<ch>, <ng>, <th>, <sch> usw.), obschon sie nur für einen Laut stehen.
- b) Die Sprachlaute, die sie repräsentieren, unterscheiden sich manchmal nur durch ein bestimmtes Merkmal voneinander (z. B. durch das Merkmal Stimmhaftigkeit im Fall von /z/ - <s> und /s/ - <ß>).
- c) Durch dialektale Besonderheiten können diese z. T. minimalen Unterschiede verwischt werden (z. B. <ch> und <sch> im „mitteldeutschen Streifen“).
- d) Einige spezielle Grapheme sind darüber hinaus sehr häufig (z. B. das Graphem <f>), andere dagegen relativ selten (z. B. das Graphem <v>).
- e) Außerdem sind häufige und seltenere Grapheme z. T. einander paarig zugeordnet (z. B. <f> und <v>) und sorgen durch ihre unterschiedlichen Vorkommenshäufigkeiten für zusätzliche Verwirrung.

Auch andere Diagnoseverfahren widmen sich diesen oder ähnlichen orthographischen Phänomenen, wenn auch mit unterschiedlichen Systematiken. Die HSP (vgl. May, 2013: 15) führt „besondere Grapheme“ wie <x>, <qu>, <v>, <ß> auf. Die OLFA klassifiziert Verwechslungen zwischen <e> und <ä>, <f> und <v>, <w> und <v>, <ch> und <g> sowie bestimmte Fremdwortfehler als eigenständige Fehlerkategorien (vgl. Thomé & Thomé, 2020: 23 ff.).

Auf Basis der AFRA-Systematik konzentriert sich die Annotationsebene [Graph_SG] im DeChiLKo auf die Verwechslungen zwischen <e> und <ä>, <sch> und <s>, <f> und <v> sowie <ch> und <g> am Endrand. Fehlschreibungen bei sonstigen speziellen Graphemen (z. B. <ng>, <th>) werden als SG° markiert. Darüber hinaus werden fehlerhafte Graphemwahl in Fremdwörtern (z. B. *<K> statt <C> im Wort *Café*) sowie transferbedingte Abweichungen (z. B. <c> für <k> in *<communication>) als FWG (Fremdwortgraphem) annotiert²¹. Da Abweichungen bei den Graphemen <z>, <s> und <ß> bereits ausführlich unter [Graph_S-Schreibung] behandelt wurden (siehe 4.2.4), stehen diese Abweichungen hier nicht im Fokus.

Insgesamt wurden 318 Abweichungen in der Kategorie [Graph_SG] in 162 Lerner-texten gefunden, was einem Fehlerwert von 10,04 pro 1.000 Tokens entspricht (vgl. Tabelle 25).

20 Aus systemlinguistischer Perspektive ist Kriterium b) nicht uneingeschränkt haltbar: Auch Laute, die sich nur durch ein einziges phonologisches Merkmal unterscheiden (z. B. [j] vs. [y] durch das Merkmal Lippenrundung), führen nicht automatisch zu orthografischen Schwierigkeiten. Die Klassifikation der AFRA beruht in diesem Punkt daher eher auf empirisch-didaktischen Beobachtungen als auf einem rein systemlinguistischen Kriterium.

21 Die gesonderte Annotation als FWG (statt als SG°) begründet sich aus der Fragestellung dieser Arbeit, die u. a. den Einfluss des mehrsprachigen Repertoires der Lernenden untersucht. Fehlschreibungen in Fremdwörtern können zumindest teilweise auf Transfer aus anderen Sprachen des Repertoires zurückzuführen sein, weshalb ihre analytische Trennung von rein graphematisch bedingten SG°-Fehlern sinnvoll erscheint. Eine eindeutige Zuordnung der Ursachen ist im Einzelfall jedoch nicht immer möglich.

Tabelle 25: Überblick über die Häufigkeit der SG-Abweichungen

SG-Abweichungen insgesamt	Summe aller Tokens	SG-Annotationen pro 1.000 Tokens
318	31.674	10,04

Die Verteilung der SG-Abweichungen nach Fehlerarten (vgl. Abbildung 21) zeigt, dass die Verwechslung zwischen <f> und <v> (FVS) mit einem Fehlerwert von 3,13/1.000 Tokens die häufigste Fehlerquelle darstellt. An zweiter Stelle stehen Fehler bei sonstigen speziellen Graphemen (SG*) mit 2,65/1.000 Tokens. Mit deutlichem Abstand folgen Verwechslungen zwischen <sch> und <s> (SCHS) (1,11/1.000 Tokens), <e> und <ä> (EÄS) (1,04/1.000 Tokens) und Fehler bei Fremdwortgraphemen (FWG) (1,01/1.000 Tokens). Die Verwechslungen zwischen <w> und <v> (WVS) sowie <ch> und <g> (CHGS) kommen wenig im Korpus vor.

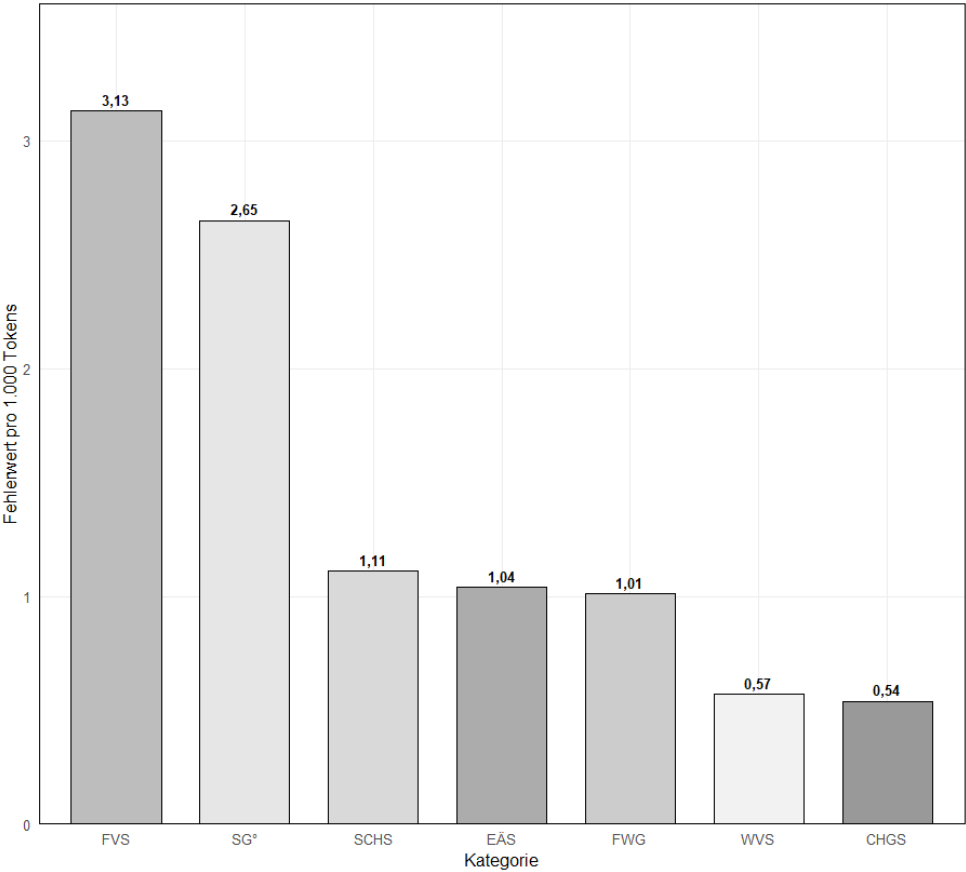


Abbildung 21: Fehlerverteilung der speziellen Grapheme nach Kategorien; Fehlerwerte pro 1.000 Tokens

Die detailliertere Betrachtung der spezifischen Fehlerarten in der Tabelle 26 liefert weitere Einblicke: Innerhalb der häufigsten Kategorie FVS dominiert die Fehlschreibung <f> für <v> (1,99/1.000 Tokens) deutlich gegenüber der umgekehrten Verwechslung <v> für <f> (1,14/1.000 Tokens). In der Kategorie SG° treten die meisten Fehler bei der Verschriftlichung der Phoneme auf, die durch Mehrgrapheme repräsentiert werden. Häufig sind hier die fehlerhafte Graphemwahl <tz> statt <ts> (0,66/1.000 Tokens) bei dem Laut [ts] sowie <ch> für <sch> (0,57/1.000 Tokens), wobei letztere oft das Suffix {-isch} betrifft (eine genauere Analyse der Suffixschreibung erfolgt in Kapitel 4.4.2). Eine weitere bemerkenswerte Fehlerart ist die Schreibung <n> für <ng> (0,44/1.000 Tokens). Für das Phonem /ŋ/ sind im Deutschen sowohl <ng> (z. B. *Menge*) als auch <n> (z. B. *Bank*) möglich. Nach Thomé et al. (2016) gilt <ng> aufgrund seiner häufigeren Repräsentation als Basisgraphem, während <n> in diesem Kontext als Orthographem fungiert. Obwohl Orthographeme häufig als „schwieriger“ gelten (vgl. Siekmann, 2021: 165), macht die Fehlerhäufigkeit der Schreibung <n> für <ng> deutlich, dass das Orthographem nicht zwangsläufig fehleranfällig ist, sondern hier möglicherweise eine Vereinfachungsstrategie oder eine falsche Analogiebildung vorliegt. Bei der Verwechslung zwischen <e> und <ä> (EÄS) ist die häufigste Fehlschreibung die Verwendung des Orthographems <ä> für das Basisgraphem <e> (1,04/1.000 Tokens) bei der Realisierung des Kurzvokals /ɛ/. Die Kategorie Fremdwortgrapheme (FWG) illustriert, dass die Schreibung <c> für <k> mit einem Fehlerwert von 0,19 eine relevante Fehlerquelle darstellt. Beim Laut [j] ist die Fehlschreibung <sch> für <s> häufiger als die Verwechslung in umgekehrter Richtung. Schließlich sind die fehlerhaften Graphemauswahlen <w> statt <v> beim Phonem /v/ (0,54/1.000 Tokens) sowie <ch> für <g> beim Phonem /ç/ im Silbenendrand (0,54/1.000 Tokens) zu nennen.

Tabelle 26: Fehlerverteilung der speziellen Grapheme nach Fehlerarten

Kategorie	Fehlerart	Summe	Fehlerwert pro 1.000 Tokens
FVS	<f> für <v> (auch <F> für <v>)	63	1,99
	<v> für <f> (auch <v> für <F> und <V> für <F>)	36	1,14
SG°	<tz> für <ts>	21	0,66
	<ch> für <sch>	18	0,57
	<n> für <ng>	14	0,44
	<sch> für <ch>	8	0,25
	<n> für <ng>	5	0,16
EÄS	<ä> für <e> (auch <ä> für <E>)	24	0,76
	<e> für <ä>	9	0,28

(Fortsetzung Tabelle 26)

Kategorie	Fehlerart	Summe	Fehlerwert pro 1.000 Tokens
FWG	<c> für <k> (auch <C> für <K> und <c> für <K>)	6	0,19
	<K> für <C>	5	0,16
	<i> für <y>	4	0,13
	<ee> für <é>	3	0,09
SCHS	<sch> für <s> (auch <Sch> für <S> und <sch> für <S>)	20	0,63
	<S> für <SCH> (auch <s> für <sch>)	15	0,47
WVS	<w> für <v>	11	0,35
	<v> für <w> (auch <V> für <W>)	7	0,22
CHGS	<ch> für <g>	9	0,28
	<g> für <ch>	8	0,25

Im Kontext des multikompetenten Repertoires der chinesischen Deutschlernenden lassen sich viele der im Bereich spezieller Grapheme beobachteten Abweichungen erklären, die durch Transferprozesse aus der L1 und der ersten Fremdsprache Englisch sowie durch die inhärenten Schwierigkeiten des deutschen Orthographiesystems erklärbar erscheinen.

Die dominante Fehlschreibung <f> für <v> ist ein klassischer Fehlerbereich im Orthographieerwerb (vgl. Thomé et al., 2016; Siekmann, 2021). Die deutsche Orthographie weist hier eine gewisse Komplexität auf: Obwohl das Phonem /f/ überwiegend mit <f> realisiert wird (65,38 % der Schreibungen für /f/), existiert auch die Schreibung mit <v> (31,95 %), insbesondere in häufigen Morphemen wie {ver-} und {vor-}, die allein fast 45 % der <v>-Schreibungen für /f/ ausmachen (vgl. Thomé et al., 2016). Für chinesische Deutschlernende ist die Phonem-Graphem-Zuordnung <f>-/f/ sehr vertraut, da sie sowohl in dem *Pinyin*-System als auch in der ersten Fremdsprache Englisch existiert. Die Untersuchung von Siekmann (2021) an freien Schülertexten deutscher Grundschüler (Klasse 3–5) bewies jedoch, dass Schwierigkeiten mit der f/v-Schreibung nicht allein L2-spezifisch sind. Siekmann (ebd.) differenziert Fehlschreibungen zum Phonem /f/ in Kategorien wie Basisgraphem (<f>) für Orthografem (<v>, <ff>, <ph>) und umgekehrt. Ihre Ergebnisse, wonach 77,46 % der Fehlschreibungen zu /f/ den Typ „Basisgraphem für Orthografem“ (z. B. <f> für <v>) darstellen und nur 22,11 % den Typ „Orthografem für Basisgraphem“ (z. B. <v> für <f>), zeigen eine ähnliche Fehlerverteilung wie im DeChiLKo. Dies deutet darauf hin, dass die Komplexität der deutschen f/v-Schreibung selbst bereits eine generelle Lernhürde darstellen könnte. Für die chinesischen Lernen-

den kommt erschwerend hinzu, dass das Graphem <v> in ihrer L1 nicht vorkommt. Obwohl ihnen das Graphem <v> aus dem Englischlernen prinzipiell vertraut ist, repräsentiert das deutsche Graphem <v> häufig den Laut [f] (z. B. in *Vater*). Dieses komplexe Zusammenspiel der Vorkenntnisse aus der L1 und der L2 Englisch könnte somit zu den beobachteten erhöhten Fehlerwerten bei der deutschen *f/v*-Schreibung beitragen.

Die Fehlschreibung <c> für <k> in der Kategorie FWG ist ein deutlicher Hinweis auf einen Transfereffekt aus der ersten Fremdsprache Englisch. Im Englischen ist die Graphem-Phonem-Korrespondenz <c> für /k/ (z. B. *cat*, *communication*) sehr verbreitet. Da die meisten chinesischen Lernenden Englisch vor Deutsch lernen, ist es plausibel, dass sie dieses etablierte Schreibmuster aus ihrem multikompetenten Repertoire auf deutsche Wörter übertragen, insbesondere wenn es sich um Fremdwörter handelt. Eine nähere Betrachtung der Fehler bei Fremdwortschreibung erfolgt in Kapitel 4.6.1.

Die Schwierigkeiten bei der Unterscheidung von <sch> und <s> für den Laut [ʃ] (SCHS) könnten möglicherweise mit der Allographie des deutschen Graphems <s> in Zusammenhang stehen. Das Graphem <s> kann je nach Kontext die Phoneme /s/, /z/ oder /ʃ/ repräsentieren. Andererseits existiert diese spezifische Mehrdeutigkeit im *Pinyin*-System nicht, wo der Laut [ʃ] (dem deutschen /ʃ/ phonetisch ähnlich) mit <sh> und /s/ mit <s> relativ eindeutig kodiert werden. Die im DeChiLKo häufigere Fehlschreibung <sch> für <s> im Vergleich zur umgekehrten Richtung könnte auf eine Übergeneralisierung des im Deutschen häufigsten Graphems für /ʃ/ hindeuten, insbesondere wenn die spezifischen Kontexte für die <s>-/ʃ/-Zuordnung noch nicht internalisiert sind.

Ein bemerkenswerter Befund ist die Verwechslung von <e> und <ä> (EÄS), insbesondere die häufigere Schreibung <ä> für <e> beim Kurzvokal /ɛ/. Die Lernenden tendieren hier dazu, das zielsprachenspezifische Graphem <ä> anstelle des Basisgraphems <e> zu verwenden. Dies könnte auf eine Übergeneralisierung der Umlautmarkierung oder eine Unsicherheit hinsichtlich der genauen Anwendungsbedingungen von <ä> versus <e> für den Laut /ɛ/ hindeuten.

Die Fehler bei sonstigen speziellen Graphemen (SG°), wie <tz> für <ts> oder <ch> für <sch>, deuten auf Schwierigkeiten mit der korrekten graphemischen Repräsentation spezifischer deutscher Di- bzw. Trigraphen hin. Die Fehlschreibung <n> für <ng> könnte auf Transfereffekte aus bestimmten chinesischen Heimatdialekten zurückzuführen sein. In vielen chinesischen Dialekten (u. a. Wu-, Xiang-, Gan-, Yue-Dialekt) werden die Auslaute /n/ und /ŋ/ nicht voneinander unterschieden, vgl. T. Liu (2015, S. 85–88). Für Lernende aus solchen Regionen könnten Wörter wie /lan/ und ein hypothetisches */lan/ phonetisch identisch klingen. Besonders im Diktat, wo der auditive Input eine entscheidende Rolle spielt, können sich solche dialektalen Einflüsse aus der L1 in einer Verwechslung von <n> für <ng> bei der Verschriftlichung wie z. B. *<lan> (*lang*) widerspiegeln.

Zusammenfassend lässt sich sagen, dass die Schreibung spezieller Grapheme für chinesische Deutschlernende eine notable Fehlerquelle darstellt. Die Fehler sind oft

auf eine komplexe Interaktion von L1-Interferenzen (sowohl aus dem Standardchinesischen als auch aus Dialekten), Englisch-Einflüssen und den inhärenten Komplexitäten des deutschen Orthographiesystems zurückzuführen.

4.3 Der silbische Bereich

Im silbischen Bereich der Orthographie werden Abweichungen im Hinblick auf ihre Position innerhalb der graphematischen Silbenstruktur untersucht. Die graphematische Silbe fungiert hierbei als eine strukturelle Einheit, die zwischen dem einzelnen Graphem und dem graphematischen Wort angesiedelt ist. Fehler werden entsprechend ihrer Lokalisation am Anfangsrand, im Silbenkern oder am Endrand der graphematischen Silbe markiert. Gemäß der Definition von Fuhrhop und Peters (2013, 216 f.) besteht der Kern einer graphematischen Silbe aus einem oder mehreren kompakten Buchstaben, die typischerweise phonologische Vokale repräsentieren. Die Ränder hingegen bestehen aus nicht kompakten Buchstaben (Konsonantenbuchstaben), deren Abfolge innerhalb des Randes spezifischen graphotaktischen Regeln folgt.

Obwohl zwischen phonologischen und graphematischen Silben viele Verbindungen und Parallelitäten bestehen, gibt es jedoch Unterschiede. So muss beispielsweise jede graphematische Silbe einen kompakten Kern haben, während phonologisch nicht jede Silbe zwingend einen Vokal enthalten muss. Des Weiteren können graphematische Silben mit einem kompakten Kern beginnen. Dies stellt in der Phonologie eine Abweichung dar, weil vokalisch anlautende, betonte Silben im Deutschen in der Regel mit einem glottalen Verschlusslaut [ʔ] als konsonantischem Anfangsrand realisiert werden (z. B. <arm> /ʔaʁm/). Die genaue phonologische Klassifikation dieses glottalen Anlauts ist dabei umstritten. Ein weiterer Unterschied betrifft die Silbifizierung mehrsilbiger Wörter: In der Graphematik wird die zweite Silbe im Vergleich zur Phonologie oft als „bedeckt“ betrachtet, wie beispielsweise im Wort <Ruhe>. Hier legt das silbeninitiale <h> eine Segmentierung nahe, die einen nicht kompakten Buchstaben am Anfangsrand der zweiten Silbe postuliert, obwohl phonologisch an dieser Stelle kein Konsonant realisiert wird²² (vgl. Fuhrhop & Peters, 2013, S. 228).

Im Rahmen der vorliegenden Untersuchung werden im DeChiLKo-Korpus Fehler im silbischen Bereich mithilfe folgender Annotationsebenen erfasst:

- [Silben_AR]: Abweichungen am Anfangsrand der Silbe (Unterscheidung zwischen einfachem und komplexem Anfangsrand);
- [Silben_SK]: Fehler im Silbenkern, klassifiziert nach der Art des betroffenen Vokals bzw. Vokalgraphems (z. B. Monophthong, Diphthong, Schwa-Laut);
- [Silben_ER]: Fehler am Endrand der Silbe (Unterscheidung zwischen einfachem und komplexem Endrand).

22 Empirische Aussprachedaten (AADG; Kleiner, 2011 ff.) belegen jedoch, dass eine tatsächliche artikulatorische Realisierung des /h/ als Hauchlaut in solchen silbeninitialen Positionen im aktuellen Gebrauchsstandard durchaus vorkommen kann (Kleiner, 2011 ff.). Für die hier diskutierte graphematische Systematik ist jedoch primär die phonologische Abstraktion der Standardlautung anzusetzen.

Die Ebenen [Silben_Fehler] und [Silben_ZH] dienen dabei der grundlegenden Segmentierung des Lernertexts bzw. der Zielhypothese nach Silbenstrukturen.

Aufgrund technischer Beschränkungen bei der Segmentierung in EXMARaLDA werden Abweichungen bei der Silbengelenkschreibung dem Endrand der ersten beteiligten Silbe zugeordnet, im Bewusstsein, dass ein Silbengelenk definitionsgemäß eine Brücke zwischen zwei Silben bildet. Fehlerhafte Schreibungen des silbeninitialen <h> werden dem Anfangsrand der betreffenden Silbe zugeordnet.

4.3.1 Fehler am Anfangsrand [Silben_AR]

Der Anfangsrand einer Silbe umfasst sämtliche Konsonantengrapheme, die dem Silbenkern innerhalb derselben Silbe vorausgehen. Im Deutschen kann der Anfangsrand einer Silbe aus bis zu drei Konsonanten bestehen oder aber leer sein (vgl. Fuhrhop & Peters, 2013, 216 f.). Besteht der Anfangsrand nur aus einem Konsonanten, wird er als einfacher Anfangsrand bezeichnet, während ein komplexer Anfangsrand mehr als zwei Konsonanten aufweist.

In der Annotationsebene [Silben_AR] werden Abweichungen im einfachen Anfangsrand als fehlerhaft (*ein), fehlend (ein-) oder überflüssig (ein+) markiert. Abweichungen beim komplexen Anfangsrand werden entsprechend als fehlerhaft (*kom), fehlend (kom-) oder überflüssig (kom+) annotiert. Das silbeninitiale <h>, das als nicht kompakter Buchstabe fungiert, wird ebenfalls im Rahmen dieser Annotationsebene berücksichtigt und als spezifischer Fehlertyp (SinH) erfasst, wenn es fehlerhaft gesetzt oder ausgelassen wird.

Im gesamten DeChiLKo wurden 839 Abweichungen am Anfangsrand in 240 Lernertexten identifiziert. Dies entspricht einem Fehlerwert von 26,49 pro 1.000 Tokens (siehe Tabelle 27).

Tabelle 27: Überblick über die Abweichungen am Anfangsrand

Abweichungen am Anfangsrand insgesamt	Summe aller Tokens	Fehler am Anfangsrand pro 1.000 Tokens
839	31.674	26,49

Das Fehlerverteilungsdiagramm (vgl. Abbildung 22) deutet darauf hin, dass die Fehler am Anfangsrand vornehmlich bei den einfachen Anfangsrändern auftreten. Die häufigste im Korpus vorkommende Fehlerart ist der falsche einfache Anfangsrand (*ein) mit einem Fehlerwert von 16,20 pro 1.000 Tokens (wie beispielsweise die Ersetzung von <v> durch <f> in *<füreinsamen> statt <vereinsamen>). Mit deutlichem Abstand folgt der fehlende einfache Anfangsrand (ein-, z. B. *<eute> statt <heute>) mit einem Fehlerwert von 3,85 pro 1.000 Tokens. Auf dem dritten und vierten Rang finden sich der falsche komplexe Anfangsrand (*kom, z. B. <st> statt <str> in *Straße*) sowie der überflüssige einfache Anfangsrand (ein+, z. B. *<halles> statt <alles>). Weniger häufig sind Abweichungen bei der Schreibung des silbeninitialen <h> (SinH, z. B. *<ruen> statt <ruhen>) sowie fehlende komplexe Anfangsränder (kom-). Die Fehlerart „überflüssiger komplexer Anfangsrand“ (kom+) tritt im Korpus nicht auf.

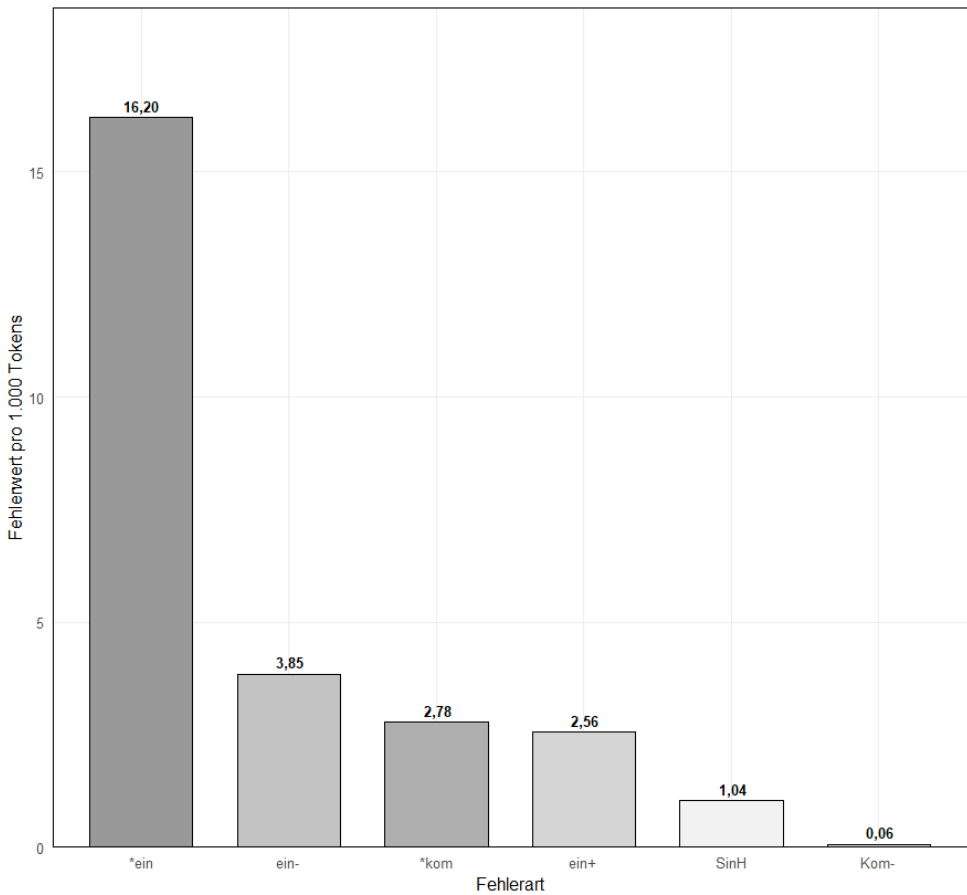


Abbildung 22: Fehlerverteilungen am Anfangsrand nach Fehlerarten; Fehlerwert pro 1.000 Tokens

Obwohl der phonotaktische Vergleich zwischen dem Deutschen und dem Chinesischen vermuten lässt, dass komplexe Anfangsränder im Deutschen (z. B. <st> in *Stadt* (/ʃtat/) oder <str> in *Straße* (/ˈʃtʁa:sə/) für chinesische Deutschlernende eine besondere Herausforderung darstellen könnten – da im Chinesischen am Silbenanfang, im Gegensatz zu den vielfältigen Möglichkeiten im Deutschen, maximal nur ein Konsonant vorkommt –, belegt die Fehlerverteilung bei Anfangsrändern im DeChiLKo, dass die Lernenden häufiger Fehler bei einfachen Anfangsrändern machen. Als primäre Ursache hierfür kann zum einen die generell höhere Text- und Silbenfrequenz einfacher Anfangsränder im Deutschen angeführt werden, wodurch die Lernenden im Input quantitativ weitaus häufiger mit diesen Strukturen konfrontiert werden. Andererseits ist es denkbar, dass bewusst mehr Aufmerksamkeit beim Erwerb von sprachlichen Unterschieden zwischen Erst- und Zielsprache den als komplexer erkannten Strukturen, wie eben den komplexen Anfangsrändern, gewidmet wird, was zu einer geringeren Fehlerrate in diesem spezifischen Bereich führen könnte (vgl. R. Ellis, 1997, S. 660).

Nach dieser allgemeinen Betrachtung der Fehlerverteilung wird der Fokus nun auf die detaillierte qualitative und quantitative Analyse der häufigsten Fehlerarten innerhalb der Kategorien „falscher einfacher Anfangsrand“, „fehlender einfacher Anfangsrand“ sowie „falscher komplexer Anfangsrand“ gerichtet (siehe Tabelle 28).

Tabelle 28: Typische Fehlerarten am Anfangsrand

Kategorie	Fehlertypen	Summe	Fehlerwert pro 1.000 Tokens
*ein	<f> für <v>	62	1,96
	<mm> für <m>	37	1,17
	<v> für <f>	30	0,95
	<l> für <n>	20	0,63
	<r> für <h>	19	0,60
ein-	<ø> für <n>	22	0,69
	<ø> für <v>	19	0,60
	<ø> für <t>	17	0,54
	<ø> für <h>	14	0,44
*kom	<st> für <str>	9	0,28
	<n> für <gn>	5	0,16

Unter der Kategorie „falscher einfacher Anfangsrand“ (*ein) gilt die Verwechslung <f> für <v> in dem Präfix {ver-} als die häufigste Fehlerart (1,96/1.000 Tokens). Hierbei wird häufig die Präposition *für* anstelle des Präfixes geschrieben. Eine weitere häufige Fehlerart ist die Ersetzung von <m> durch <mm> (1,17/1.000 Tokens), wie beispielsweise im Wort *<vereinsammen> (*vereinsamen*), was auf Schwierigkeiten bei der korrekten Anwendung der Konsonantenverdopplungsregeln hindeutet, die hier jedoch am Anfangsrand einer Folgesilbe auftritt. Auch die Substitution von <f> durch <v> (0,95/1.000 Tokens) ist auffällig. Da die *f/v*-Schreibung und die Konsonantenverdopplung bereits in den Abschnitten 4.1.5 bzw. 4.1.3 thematisiert wurden, konzentriert sich die weitere Diskussion hier auf andere spezifische Fehler im einfachen Anfangsrand.

Die Substitution von <n> durch <l> am Anfangsrand weist auf eine häufige Verwechslung von Lauten [n] und [l] auf. Für chinesische Lernende ist die Unterscheidung dieser beiden Laute auch in ihrer L1 nicht immer unproblematisch. Obwohl das Hochchinesische eine klare phonetische Unterscheidung zwischen [n] und [l] kennt, werden diese Laute im Anlaut in vielen südchinesischen Dialekten (z. B. Gan-, Yue-, Wu-Dialekt) nicht voneinander unterschieden (vgl. T. Liu, 2015, 85 ff.). Solche dialektalen phonetischen Merkmale der L1 können die deutsche Aussprache beeinflussen und sich auf der schriftlichen Ebene in der Substitution <n> durch <l> am Anfangsrand manifestieren.

Auch die Ersetzung von <r> durch <h> (0,60/1.000 Tokens) am Anfangsrand lässt sich als typischer Interferenzfehler aus der L1 Chinesisch erklären. Die Phoneme /h/ und /r/ existieren in dieser Form nicht im Mandarin-Chinesischen. In Studien von C. Liu (1982), Spiekermann (1997) und Hachenberg (2003) wird bestätigt, dass das deutsche /r/ in der Aussprache chinesischer Deutschlerner häufig durch ein stark reibendes /x/ (Pinyin: <h>) ersetzt wird. Auf der schriftlichen Ebene kann diese phonetische Substitution zu Verwechslungen von <r> mit <h> führen.

In der Kategorie „fehlender einfacher Anfangsrand“ (ein-) tritt das Auslassen des Anfangsrandes besonders häufig in mehrsilbigen Wörtern auf und betrifft dabei größtenteils nicht den Anfangsrand der ersten Wortsilbe, sondern eher den der nachfolgenden Silben. Beispiele hierfür sind *<meien> (*meinen*), *<mutiert> (*motiviert*) oder *<Weiterlernen> (*Weiterlernen*). Besonders auffällig ist das Fehlen von <n> am Anfangsrand bei den Wörtern *Belohnungen* (n = 13) und *ablehnen* (n = 3), wobei das Graphem <n> nach einem mit Dehnungs-<h> markierten Langvokal (<oh> (/o:/) bzw. <eh> (/e:/)) ausgelassen wird. Die korrekte Schreibung des vorausgehenden Langvokals deutet darauf hin, dass die Lernenden die Markierung der Vokallänge mittels Dehnungs-<h> prinzipiell erworben haben. Die Schwierigkeit mit Konsonantenclustern bei chinesischen Deutschlernenden ist bereits in vielen Studien dokumentiert worden (vgl. Hunold, 2009, S. 168; Qiu, 1984a, S. 126). Die vermehrte Auslassung des Graphems <n> nach einem Dehnungs-<h> (<oh-n> oder <eh-n>) könnte somit als eine Strategie zur Vermeidung einer als komplex empfundenen Graphemhäufung interpretiert werden, auch wenn es sich hier nicht um ein klassisches Konsonantencluster im phonetischen Sinne handelt, sondern um eine graphematische Sequenz.

Im Gegensatz zu den anderen Auslassungen in dieser Kategorie betrifft die Weglassung des einfachen Anfangsrandes <v> hauptsächlich den Wortanfang, und zwar spezifisch das Präfix {ver-}. Da hier oft das gesamte Präfix betroffen ist, wird diese Fehlerart im Rahmen der morphologischen Fehleranalyse in Kapitel 4.4.2 genauer betrachtet.

Die Auslassung des einfachen Anfangsrandes <r> tritt gehäuft (n = 7) im Wort *Fahrrad* auf, wobei statt <Fahrrad> häufig *<Fahrt>²³ geschrieben wird, was auf eine Assimilation oder Elision des ersten <r> aufgrund der nachfolgenden identischen Laut- bzw. Graphemsequenz hindeutet. Ähnliche Auslassungen finden sich auch in <*überachent> (*überraschend*). Solche Phänomene werden in der HSP als „Phonemverschmelzung“ klassifiziert und als Verstöße gegen die morphologische Strategie eingeordnet (vgl. May, 2013, S. 35); eine detailliertere Betrachtung erfolgt daher im morphologischen Teil dieser Arbeit.

In der Fehlerkategorie „falscher komplexer Anfangsrand“ (*kom) ist die häufigste Abweichung die Reduktion von <str> zu <st> (0,28/1.000 Tokens). Im deutschen Konsonantencluster /ʃʏʁ/ treffen zwei für chinesische Lerner fremde Konsonanten (/ʃ/ und /ʏʁ/) aufeinander. Die Kombination aus Konsonantenhäufung und fremden Lauten stellt eine besondere Lernschwierigkeit dar. Die Vereinfachung zu /ʃt/ ist eine typische Reduktionsstrategie.

23 Auch *<fahrt>, *<fat>, *<fährt> usw.

Die Verwechslung von <n> statt <gn> (0,16/1.000 Tokens) ist ein weiteres Beispiel für die Schwierigkeiten mit komplexen Konsonantenclustern. Die korrekte Aussprache des Clusters /gn/ (<gn>) erfordert eine schnelle Umstellung der Zungenartikulation vom velaren /g/ zum alveolaren /n/, was für chinesische Lernende aufgrund fehlender ähnlicher Strukturen in der L1 anspruchsvoll sein kann. Bei der Konfrontation mit komplexen Clustern tendieren Lernende häufig zur Reduktion. Die Reduktion von Konsonantenverbindungen ist einerseits ein typischer phonologischer Prozess, der auch im L1-Erwerb des Deutschen beobachtet wird (vgl. Fox-Boyer, Groos & Schauß-Golecki, 2015, S. 28). Andererseits neigen Lernende, deren L1 nur stark vereinfachte oder keine Konsonantenverbindungen aufweist, im L2-Erwerb des Deutschen ebenfalls dazu, komplexe Verbindungen zu reduzieren (vgl. Flege, 1995, S. 241). Dieser Reduktionsprozess kann sich in der Schreibung als Fehlschreibung <n> statt <gn> manifestieren.

Zusammenfassend lässt sich für den Bereich des Anfangsrandes festhalten, dass Fehler häufiger bei einfachen Anfangsrändern als bei den komplexen Rändern auftreten, obwohl die L1-Struktur des Chinesischen (maximal ein Konsonant im Silbenanlaut) eher Schwierigkeiten bei komplexen deutschen Anfangsrändern erwarten ließe. Dies könnte auf die höhere Frequenz einfacher Anfangsränder im Deutschen oder/und eine stärkere Lernaufmerksamkeit für die offensichtlich komplexeren Strukturen zurückzuführen sein.

Die qualitative Analyse der spezifischen Fehler im einfachen Anfangsrand verdeutlicht die Interaktion im multikompetenten System der Lernenden: Verwechslungen wie <f> für <v>, <n> für <l> und <r> für <h> lassen sich plausibel auf L1-Einflüsse (Fehlen bestimmter Phoneme im L1, dialektale Aussprachevarianten) und die Komplexität des deutschen Systems zurückführen. Das Fehlen einfacher Anfangsränder, insbesondere in nicht-initialen Silben und nach Dehnungs-<h>, deutet auf Strategien zur Vermeidung von als komplex empfundenen graphematischen Sequenzen oder auf Assimilationsprozesse hin. Fehler bei komplexen Anfangsrändern, wie die Reduktion von <str> zu <st> oder <gn> zu <n>, bestätigen die erwarteten Schwierigkeiten mit deutschen Konsonantenclustern aufgrund der L1-Phonotaktik.

Insgesamt illustrieren die Fehler am Anfangsrand, dass nicht nur die offensichtlichen strukturellen Unterschiede zwischen L1 und Zielsprache Deutsch eine Rolle spielen, sondern auch subtilere L1-Einflüsse und die inhärenten Komplexitäten des deutschen Silbenbaus für die Lernenden relevant sind.

4.3.2 Fehler im Silbenkern [Silben_SK]

Auf der Annotationsebene [Silben_SK] werden Abweichungen im Silbenkern erfasst und nach der jeweiligen Fehlerart klassifiziert. Der Art des Anschlusses zwischen dem vokalischen Element und dem nachfolgenden konsonantischen Material spielt für die deutsche Orthographie eine entscheidende Rolle (vgl. Thelen, 2010, S. 36), da sie oft die Vokalquantität und deren graphemische Markierung beeinflusst. Daher erfolgt in Anlehnung an Thelen (vgl. Thelen, 2010, S. 36) zunächst eine grundlegende Einteilung der Silbenkerne in solche mit festem Anschluss (Annotationstag „FA“) und solche mit losem Anschluss.

Silbenkerne mit festem Anschluss werden im Rahmen dieser Annotationsebene nicht weiter unterteilt, da die Fälle, die zu unterschiedlicher orthographischer Markierung führen, bereits im phonographischen Bereich unter [Graph_PGK] und [Graph_VQM] (siehe Kapitel 4.2) detailliert behandelt wurden. Bei Silbenkernen mit losem Anschluss hingegen werden die Abweichungen in drei Unterkategorien eingeteilt: Fehler bei den Monophthongen, bei Diphthongen und bei zentralisierenden Diphthongen. Letztere bezeichnen Kombinationen aus einem Vokal und vokalisiertem <r> (z. B. /i:ɐ/ in *wir*). Darüber hinaus werden die Abweichungen bei den Phonemen /ə/ und /ɐ/ in unbetonten Silben als *Schwa* markiert. Für jede dieser Fehlerarten wird zwischen falschen (*), ausgelassenen (-) und überflüssigen (+) Elementen unterschieden.

Im gesamten DeChiLKo wurden 2.669 Abweichungen im Silbenkern aus 314 Lernertexten identifiziert. Dies entspricht einem Fehlerwert von 84,26 pro 1.000 Tokens.

Tabelle 29: Überblick über Abweichungen im Silbenkern

Abweichungen im Silbenkern insgesamt	Summe aller Tokens	Abweichungen im Silbenkern pro 1.000 Tokens
2.669	31.674	84,26

Die Fehlerverteilung im Silbenkern nach den oben genannten Kategorien (siehe Abbildung 23) deutet darauf hin, dass der hohe Fehlerwert bei der Verschriftlichung von Schwa-Lauten in unbetonten Silben am auffälligsten ist. Diese Fehlerart stellt mit einem Fehlerwert von 40,89 pro 1.000 Tokens die mit Abstand größte Fehlerquelle dar. Mit deutlichem Abstand folgen Abweichungen in Silbenkernen mit Festanschluss (FA), die einen Fehlerwert von 20,52 pro 1.000 Tokens aufweisen. Ebenfalls bemerkenswert ist die Kategorie der zentralisierenden Diphthonge (ZDiph), deren Fehlerwert bei 11,81 pro 1.000 Tokens liegt. Abweichungen bei den Monophthongen mit losem Anschluss (Mono) treten mit 9,00 pro 1.000 Tokens etwas seltener auf, während die Diphthonge (Diph) mit nur 2,05 pro 1.000 Tokens die geringste Fehlerquelle im Silbenkern darstellen.

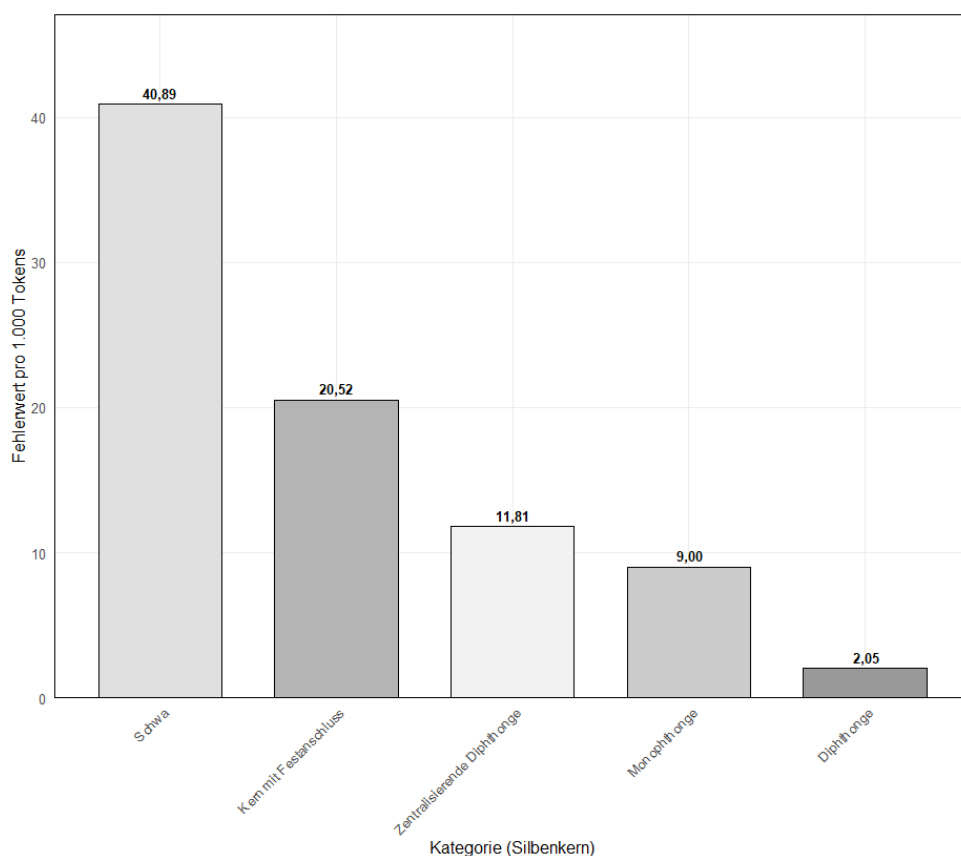


Abbildung 23: Fehlerverteilung im Silbenkern nach Kategorien; Fehlerwert pro 1.000 Tokens

Der relativ geringe Fehlerwert bei der Verschriftlichung von Diphthongen (/aʊ/, /aɪ/, /ɔɪ/) im Deutschen lässt sich plausibel durch die phonetische Struktur der L1 Chinesisch erklären. Während der deutsche Kernwortschatz lediglich diese drei Diphthonge umfasst, weist das Mandarin-Chinesische eine deutlich größere Bandbreite und Frequenz an Vokalverbindungen und Diphthongen auf (vgl. Hunold, 2009, S. 89; T. Liu, 2015, S. 38; Tang, 2003, S. 127). Dazu gehören vier fallende Diphthonge: /ai/ (<ai>), /au/ (<ao>), /ei/ (<ei>) sowie diphthongische Einheiten mit Übergangsvokal /i/: /ia/ (<ia>), /ie/ (<ie>), /iə/ (<ie>), /uo/ (<uo>), /ua/ (<ua>) und /yɛ/ (<üe>). Aufgrund dieser im multikompetenten Repertoire bereits vorhandenen Vertrautheit mit komplexen Vokalverbindungen sind chinesische Deutschlernende mit der Schreibung von Diphthongen im Deutschen vergleichsweise sicher, was der niedrigere Fehlerwert in diesem Bereich bestätigt. Im Gegensatz dazu bereitet die Schreibung der zentralisierenden Diphthonge (Kombinationen aus Vokal und vokalisiertem <r>) den Lernenden größere Schwierigkeiten.

Durch die ebenenübergreifende Analyse mit den Annotationsebenen der Graphem-Phonem-Korrespondenz [Graph_PGK] und Vokalquantitätsmarkierung [Graph_VQM] lässt sich erkennen, dass ein Großteil der Fehler (68,9%) die Auswahl einfacher Vokalgrapheme betrifft, während Fehler bei der Markierung der Vokalquantität 31,1% ausmachen. Obwohl die Monophthonge /a:/, /e:/, /o:/, /u:/, /ø:/, /y:/ häufig als „unproblematische Fälle“ bezeichnet werden (vgl. Thelen, 2010, S. 36), haben die chinesischen Lernenden vor allem Schwierigkeiten mit den Graphemen <oh>, <ie>, <o>, <ü>, <e>. Ähnlich verhält es sich bei Silbenkernen mit festem Anschluss (FA): 86,1% der Abweichungen betreffen einfache Vokalgrapheme, während 13,9% auf Fehler bei der Vokalquantitätsmarkierung zurückzuführen sind.

Die mit Abstand höchste Fehlerhäufigkeit im Bereich der Schwa-Laute ([ə] und [ɐ]) in Reduktionssilben bestätigt die bereits in Abschnitt 4.1.2 (PGK-Fehler) diskutierten Schwierigkeiten chinesischer Deutschlernender mit diesen spezifischen Schwa-Lauten, für die es im Chinesischen keine direkten Entsprechungen gibt. In der folgenden Analyse liegt der Schwerpunkt verstärkt auf den silbenstrukturellen Verhältnissen bezüglich des nachfolgenden Endrands.

Die folgende Abbildung 24 zeigt eine detaillierte Fehlerverteilung im Silbenkern nach Fehlerarten.

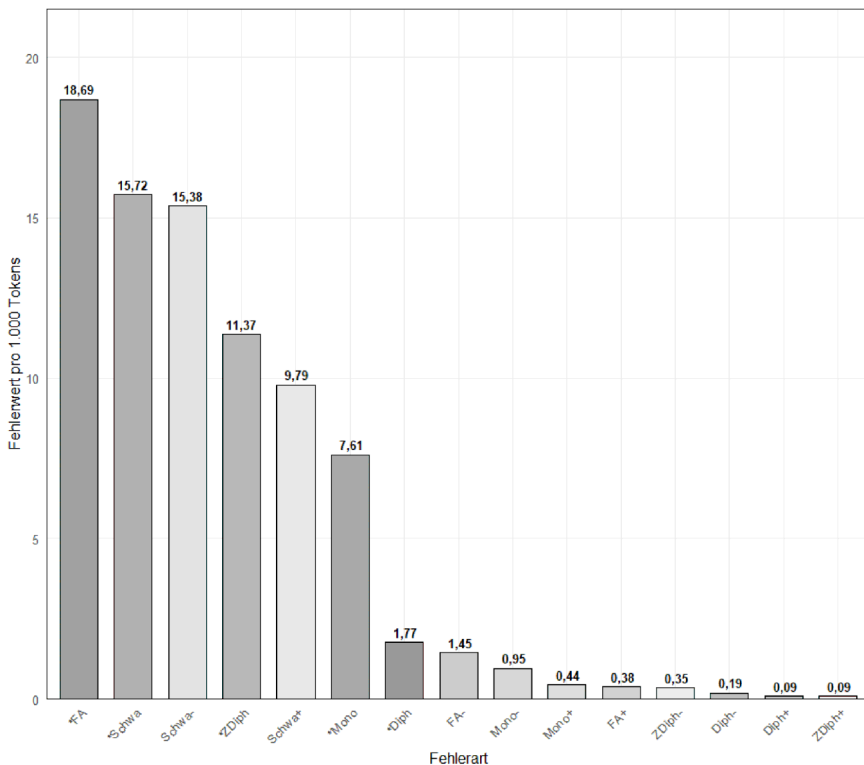


Abbildung 24: Fehlerverteilung im Silbenkern nach Fehlerarten; Fehlerwert pro 1.000 Tokens

Es ist deutlich zu erkennen, dass die Ersetzungsabweichungen innerhalb jeder Kategorie am häufigsten auftreten. Dabei können der korrekte Silbenkern und das fehlerhaft verwendete Element zumeist einer abstrakten Klasse zugeordnet werden, d. h. Monophthonge werden bevorzugt durch Monophthonge, Diphthonge durch Diphthonge und Schwa-Laute durch Schwa-Laute ersetzt. Dies kann als Hinweis darauf gewertet werden, dass die Silbe als Planungseinheit beim Schreiben aktiv ist (vgl. Weingarten, 2004, S. 10). Tabelle 30 listet die häufigsten Fehlertypen in den jeweiligen Kategorien auf.

Tabelle 30: typische Fehlertypen in den Silbenkernen

Kategorie	Fehlertypen	Summe	Fehlerwert pro 1.000 Tokens
*FA	<ie> für <e>	44	13.89
	<ei> für <e>	44	13.89
	<u> für <ü>	42	13.26
	<a> für <ei>	24	7.58
*Schwa	<e> für <er>	332	104.83
	<er> für <e>	108	34.10
*ZDiph	<o> für <or>	56	17.68
	<ür> für <er>	47	14.84
	<A> für <Ar> (<a> für <ar>)	35	11.05
*Mono	<o> für <oh>	28	8.84
	<i> für <ie>	23	7.26

Unter den Fehlern bei Silbenkernen mit festem Anschluss (*FA) ist die Fehlschreibung <ie> für <e> die häufigste Fehlerart. Danach folgen die Verwechslungen <u> für <ü> sowie <ei> für <e>. Unter den „falschen lose angeschlossenen Monophthongen“ (*Mono) treten vor allem die Abweichungen bei der Markierung von Langvokalen wie z. B. <o> für <oh> und <i> für <ie> auf. Aus dieser Betrachtung lässt sich feststellen, dass der feste oder lose Anschluss bei der Fehlerverteilung als Indiz für den Spracherwerb eine Rolle spielt. Bei Silbenkernen mit losem Anschluss haben die Lernenden eher Schwierigkeiten bei der Markierung der Vokalquantität, während sie bei Silbenkernen mit festem Anschluss eher Probleme bei der Markierung der Vokalqualität aufweisen. Eine weitere Analyse hinsichtlich der nachfolgenden Endränder legt nahe, dass ein Drittel der falschen Markierungen von Langvokalen in Silben mit festem Anschluss zusammen mit Abweichungen am Endrand auftreten, wie z. B. bei *<wohnen> (*wollen*), *<kaan> (*kann*), *<gewohnen> (*gewonnen*); *<erihnen> (*erinnern*). Dieses Ergebnis deutet darauf hin, dass die Lernenden teilweise mit dem Prinzip des Längenausgleichs vertraut sind, d. h. offene Silben werden mit einem gespannten Vokal realisiert, während

Silben mit komplexem Endrand einen ungespannten Vokal aufweisen. Die Schwierigkeiten bei der korrekten Anwendung dieses Prinzips hängen vor allem mit der Wahrnehmung und Unterscheidung der Vokalquantität zusammen. Dies stellt eine besondere Herausforderung für Lernende dar, in deren L1 keine phonologisch relevante Unterscheidung zwischen Kurz- und Langvokalen existiert (vgl. Becker & Siekmann, 2012; Ruppert & Hanulíková, 2022).

Unter der Kategorie „fehlerhaftes Schwa“ (*Schwa) sind Verwechslungen zwischen den beiden Schwa-Lauten am häufigsten, was darauf hindeutet, dass es den Lernenden auf einer Repräsentationsebene bewusst ist, an welcher Stelle ein Schwa-Laut stehen muss, auch wenn der konkrete Silbenkern noch nicht spezifiziert ist. Im Gegensatz dazu signalisieren die Fehlererscheinungen in der Kategorie „falsche zentralisierende Diphthonge“ (*ZDiph), dass chinesische Lernende mit diesem Phänomen noch Schwierigkeiten haben. In den Abweichungen <o> für <or> und <A> für <Ar> (auch <a> für <ar>) wird das vokalisierte <r> ausgelassen. Die Fehlerart <a> für <ar> gilt nach Thelen (vgl. 2010, S. 37) als ein besonderer öffnender Diphthong, weil keine weitere Öffnung ausgehend von /a/ mehr möglich ist und es daher möglich ist, dass /aʊ/ als /a:/ realisiert wird. Aus der schriftlichen Perspektive führt diese Realisierung zur Auslassung des vokalisierten <r>. Bei der Fehlerart <ür> für <er> handelt es sich um die Verwechslung zwischen der Präposition *für* und dem Präfix {ver-}. Als die Lernenden das Präfix {ver-} beim Diktat hörten, tendierten sie dazu, die ihnen vertraute Präposition *für* aufzuschreiben.

Ausgelassene Schwa-Laute (Schwa-) kommen häufig im Korpus vor. Eine nähere Betrachtung im Hinblick auf den nachfolgenden Endrand verdeutlicht, dass eine Auslassung des Schwa-Lauts besonders häufig in den Silben geschieht, wenn die Sonoranten /n/ (n = 153), /l/ (n = 6) den Schwa-Lauten folgen. Der Grund kann darin liegen, dass anstelle von Schwa die Sonoranten /m/, /l/, /n/, /ŋ/ in den Silbenkern rücken können (vgl. Fuhrhop & Peters, 2013, S. 80; Fuhrhop, 2021, 14 f.). In IPA wird der Silbenkern der Reduktionssilbe durch einen untergesetzten Strich angezeigt, wie z. B. /leːb̥n/ (*leben*). Auf der schriftlichen Ebene könnte dies dazu führen, dass die Lernenden den Schwa-Laut im Silbenkern der Reduktionssilbe auslassen.

Die Fehlerart „überflüssiger Schwa-Laut“ (Schwa+) tritt besonders häufig in der Form der Hinzufügung des Buchstabens <e> auf, was mit 254 Erscheinungen den größten Anteil der Kategorie ausmacht. Dieses Phänomen korrespondiert mit bekannten Ausspracheproblemen chinesischer Deutschlernender, die dazu neigen, an Wort- oder Silbenenden nach Konsonanten einen zusätzlichen Schwa-Laut (/ə/) zu artikulieren (vgl. Hunold, 2009, S. 155; T. Liu, 2015, S. 79; Qiu, 1984a, S. 126). Diese interlinguale Interferenz aus der L1 Chinesisch dient oft der Auflösung von im Chinesischen unüblichen Konsonantenclustern im Auslaut oder der Anpassung an die im Chinesischen vorherrschende offene Silbenstruktur (vgl. Hunold, 2009, S. 56; M. Wang, 1993, 66 ff.). Auf der schriftlichen Ebene kann sich dieser phonetische Transfereffekt in der überflüssigen Hinzufügung des Graphems <e> manifestieren.

Im Gegensatz zu den häufigen Fehlern bei Schwa-Lauten sind ausgelassene Kerne bei Vollvokalen im DeChiLKO selten. Dies ist dadurch erklärbar, dass der Silben-

kern sowohl im Deutschen als auch im Chinesischen ein obligatorisches Element der Silbenphonotaktik darstellt (siehe 2.3.2). Die Lernenden sind daher mit diesem grundlegenden Prinzip aus ihrem multikompetenten sprachlichen Wissen vertraut.

Zusammenfassend dokumentierte die Analyse der Fehler im Silbenkern, dass der Umgang insbesondere mit den im Deutschen spezifischen Schwa-Lauten und zentralisierenden Diphthongen für chinesische Deutschlernende eine signifikante Fehlerquelle darstellt. Während Diphthonge aufgrund der Vokalkomplexität im Chinesischen weniger Probleme bereiten, führen die subtilen phonetischen Unterschiede und die spezifischen graphematischen Konventionen bei Reduktionsvokalen und vokalisiertem <r> zu einer hohen Fehleranfälligkeit.

4.3.3 Fehler am Endrand [Silben_ER]

Der Endrand umfasst sämtliche Konsonantengrapheme, die dem Kern innerhalb derselben Silbe nachfolgen. Im Deutschen kann der Endrand einer Silbe aus bis zu fünf Konsonanten bestehen oder auch leer sein. Besteht der Endrand nur aus einem Konsonanten, wird er als einfacher Endrand bezeichnet, während ein komplexer Endrand mehrere Konsonanten aufweist (vgl. Fuhrhop, 2021, S. 14).

In der Annotationsebene [Silben_ER] werden die Abweichungen je nach der Anzahl der Konsonanten am Endrand zwischen einfach (ein) und komplex (kom) unterschieden. Dabei sind die Auslassung, Hinzufügung sowie Verwechslung zusätzlich mit den Tags -, + und * markiert. Des Weiteren werden Abweichungen bei der Silbengelenkschreibung auf dieser Ebene als SiG annotiert. Hierbei handelt es sich im Rahmen dieser Annotationsebene primär um die phonologische Schärfungsschreibung, bei der sich die Schreibung direkt aus der phonologischen Struktur ableiten lässt (z. B. <mm> im Wort *stammen*). Davon abzugrenzen sind morphologische Schärfungsschreibungen (z. B. <mm> in der flektierten Wortform *stammt*), die später im morphologischen Bereich (siehe Kapitel 4.4.1) betrachtet werden.

Insgesamt wurden 1.950 Abweichungen am Endrand von 310 Lernertexten im DeChiLko gefunden, was einem Fehlerwert von 61,56 pro 1.000 Tokens entspricht. Im Vergleich zu den Abweichungen am Anfangsrand (n = 839, Fehlerwert: 26,49) treten Fehler am Endrand somit deutlich häufiger auf.

Tabelle 31: Überblick über Abweichungen am Endrand

Abweichungen am Endrand insgesamt	Summe aller Tokens	Abweichungen am Endrand pro 1.000 Tokens
1950	31.674	61,56

Die Fehlerverteilung am Endrand nach Fehlerarten (vgl. Abbildung 25) verdeutlicht, dass die Abweichungen vornehmlich bei den einfachen Endrändern geschehen. Der falsche einfache Endrand (*ein) ist hierbei mit einem Fehlerwert von 19,01 pro 1.000 Tokens die häufigste Fehlerquelle. Darauf folgt der überflüssige einfache Endrand (ein+) mit einem Fehlerwert von 16,45 pro 1.000 Tokens. Mit geringem Abstand steht die Fehlerart „fehlender einfacher Endrand“ (ein-) an dritter Stelle. Generell treten weniger Abwei-

chungen bei den komplexen Endrändern auf, wobei hier der „falsche komplexe Endrand“ (*kom) häufiger erscheint als andere Fehlschreibungen dieser Unterkategorie. Bemerkenswert ist zudem das häufige Auftreten von Abweichungen bei der Silbengelenkschreibung (SiG) mit einem Fehlerwert von 4,48 pro 1.000 Tokens.

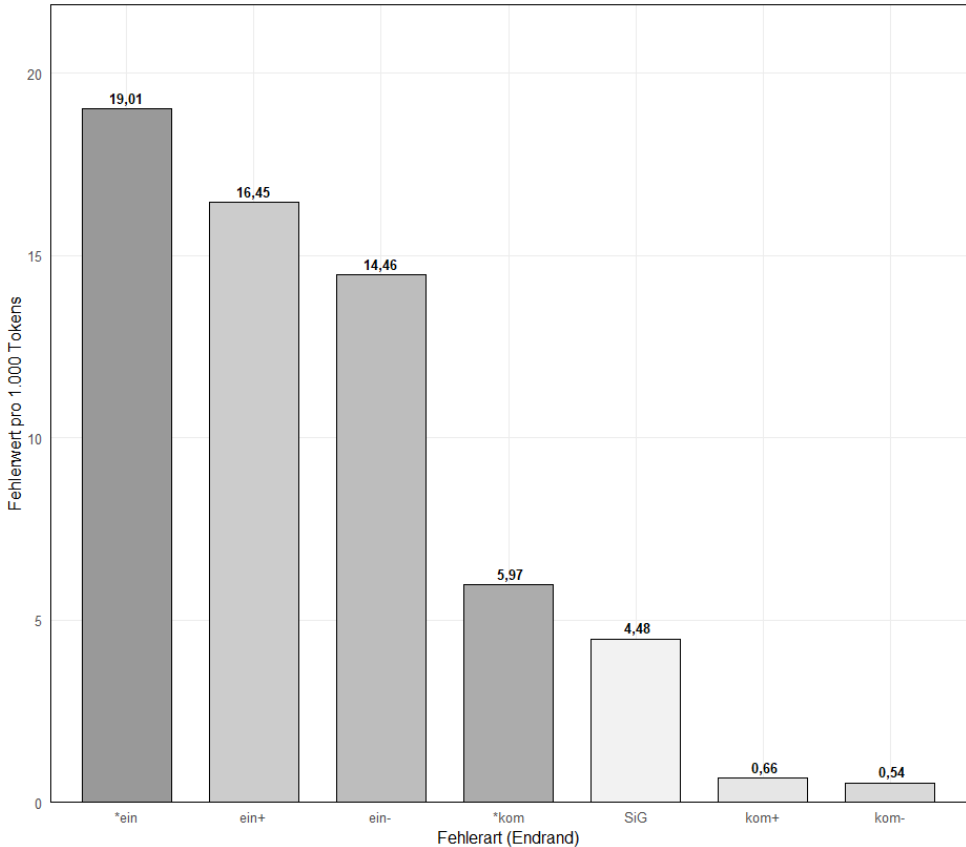


Abbildung 25: Fehlerverteilung am Endrand nach Fehlerarten; Fehlerwert pro 1.000 Tokens

Die höhere Fehleranfälligkeit des Silbenendrandes im Vergleich zum Anfangsrand bei den chinesischen Deutschlernenden könnte in der spezifischen phonotaktischen Komplexität dieser Position im Deutschen begründet sein. Im Deutschen können am Silbenendrand bis zu fünf Konsonanten aufeinandertreffen, und orthographische Regeln wie die Auslautverhärtung spielen eine wichtige Rolle. Zudem tragen Endrandkonsonanten oft grammatische Informationen (z. B. Plural, Genitiv, Tempus), was die Komplexität auf der morphosyntaktischen Ebene zusätzlich erhöht. Für chinesische Deutschlernende stellt der deutsche Silbenendrand eine besondere Herausforderung dar, weil in ihrer L1 Chinesisch eine solche grammatische Komplexität nicht vorhanden ist. Im Chinesischen endet eine Silbe typischerweise entweder auf einen Vokal oder einen Nasal

(/n/, /ŋ/), was die Silbenstruktur im Auslaut vergleichsweise einfach gestaltet (vgl. Kapitel 2.3.2).

Die Tatsache, dass Ersetzungsabweichungen (*) am Endrand – ähnlich wie am Anfangsrand und im Silbenkern – innerhalb der jeweiligen Kategorien am häufigsten auftreten, stützt erneut die Annahme, dass die Silbe beim Schreiben als Planungseinheit aktiv ist (vgl. Weingarten, 2004, S. 10). Die Lernenden scheinen die Silbenstruktur und die Notwendigkeit eines Endrandes oft zu erkennen, scheitern aber an der korrekten graphemischen Realisierung. Tabelle 32 listet die typischen Fehlertypen in den Silbenendrändern auf.

Tabelle 32: Typische Fehlertypen in den Silbenendrändern

Kategorie	Fehlertypen	Summe	Fehlerwert pro 1.000 Tokens
*ein	<n> für <m>	70	2.21
	<m> für <n>	60	1.89
	<s> für <ss>	39	1.23
	<t> für <d>	36	1.14
ein+	<n> für <ø>	381	12.03
	<l> für <ø>	85	2.68
ein-	<ø> für <n>	315	9.95
	<ø> für <l>	38	1.20
	<ø> für <m>	34	1.07
*kom	<s> für <st>	14	0.44
	<t> für <nd>	9	0.28
SiG	<s> für <ss>	13	0.41
	<ch> für <sch>	12	0.38

Unter der Kategorie „falscher einfacher Endrand“ (*ein) ist die Verwechslung zwischen den Nasalen <m> und <n> am auffälligsten. Wie bereits in Abschnitt 4.1.2 und durch die Analyse der grammatisch abzuleitenden Abweichungen ([Son_Gra-Ab]) dargelegt wurde, sind knapp 80 % dieser Verwechslungen auf die Deklinationsfehler zurückzuführen. Die ebenenübergreifende Analyse in Bezug auf die Annotationsebene „morphologische Konstanz“ [Morph_Konstanz] zeigt zudem, dass alle Ersetzungen des Endrandgraphems <d> durch <t> auf die Auslautverhärtung zurückgehen. Die Substitution von <s> für <ss> (1,23/1.000 Tokens) wird in der vorliegenden Analyse dem einfachen Endrand zugeordnet, da es sich bei der korrekten Schreibung <ss> in den untersuchten Fällen nicht primär um eine rein phonologische Silbengelenkschreibung handelt, sondern vielmehr um Fälle, die eine morphologische Betrachtung erfordern. Eine genauere Analyse dieser morphologisch relevanten Schreibphänomene erfolgt in Kapitel 4.4.1.

Besonders hervorzuheben ist die hohe Frequenz der Hinzufügung von <n> (12,03/1.000 Tokens), die die häufigste einzelne Fehlerquelle im Silbenendrand darstellt. Diese Fehlschreibung lässt sich als deutlicher Transfer aus der L1 Chinesisch erklären. Wie in Kapitel 2.3.2 dargelegt, ist der nasale Konsonant /n/ (bzw. das Graphem <n> im *Pinyin*) im Chinesischen der häufigste erlaubte konsonantische Auslaut nach dem Hauptvokal. Diese Dominanz des /n/-Auslauts in der L1 führt zu einer sogenannten Übergeneralisierung auf die Zielsprache Deutsch, sodass <n> auch dort eingefügt wird, wo es orthographisch nicht korrekt ist.

In der Kategorie „fehlender einfacher Silbenendrand“ (ein-) treten die meisten Abweichungen bei den Sonoranten (/n/, /l/, /m/) auf, wobei die Auslassung von <n> (9,95/1.000 Tokens) am auffälligsten ist. Ein wichtiger Faktor hierfür könnte die wahrnehmungsbedingte Schwächung von /n/ in der Endposition sein. Im Deutschen wird [n] im Silbenendrand oft reduziert oder nasalisiert, sodass die Lernenden das [n] im Diktat akustisch möglicherweise nicht eindeutig identifizieren und es deshalb in der Schreibung weglassen.

In der Kategorie „falscher komplexer Endrand“ (*kom) ist die häufigste Abweichung die Auslassung von <t> in der Konsonantenkombination <st> (<t> für <st>: 0,44/1.000 Tokens). Das [t] im Cluster [st] ist phonetisch weniger prominent, da es einen geringeren Sonoritätsgrad aufweist und am Wortende oft unbetont ist. Es ist daher plausibel anzunehmen, dass die Lernenden das [t] in der rezeptiven Verarbeitung des Diktats möglicherweise erschwert auditiv diskriminieren oder nicht als einen relevanten Bestandteil identifizieren, was zu dessen Auslassung in der Verschriftlichung führt. Auch die Reduktion von <nd> zu <t> (0,28/1.000 Tokens) deutet auf Vereinfachungsstrategien bei komplexen Endrändern hin.

Die Silbengelenkschreibung (SiG) wird im Korpus am häufigsten durch eine Vereinfachung der Geminata <ss> zu <s> abgebildet, wie beispielsweise in den Fehlerformen *<Wasa> (*Wasser*) oder *<sträsig> (*stressig*). Obwohl in der ersten Fremdsprache Englisch ebenfalls Konsonantenverdopplungen vorkommen, sind diese jedoch, anders als die primär phonologisch basierte Geminata im Deutschen, meist morphologisch motiviert, etwa bei der Flexion oder Ableitung. Ein Beispiel hierfür ist *running*, wo der Konsonant <n> verdoppelt wird, um die Kürze des Vokals /ʌ/ im Stammwort *run* anzuzeigen. Darüber hinaus gibt es im Englischen auch Wörter, bei denen phonetisch das Konzept eines Silbengelenks existiert, wie etwa bei <bb> in *hobby* oder <nn> in *dinner*. Die Fehlererscheinungen bei aus dem Englischen entlehnten Wörtern, wie die im Korpus beobachtete Schreibung *<hoby> (für *Hobby*), könnten ein Indiz dafür sein, dass chinesische Deutschlernende tendenziell gewisse Unsicherheiten bei der korrekten Verschriftung von Silbengelenken aufweisen. Demzufolge scheint dieses Phänomen sowohl in der ersten Fremdsprache Englisch als auch in der Zielsprache Deutsch eine potenzielle Fehlerquelle darzustellen.

Darüber hinaus ist bei der Analyse der Silbengelenkschreibung im DeChiLKo-Korpus zu beobachten, dass bei den Mehrgraphen <sch> und <ch> im Silbengelenk keine fehlerhaften Verdopplungen (z. B. *<schsch> oder *<chch>) im Silbengelenk auftreten. Dies deutet darauf hin, dass die Lernenden die Regel, dass Bi- und Trigraphen im Deut-

schen nicht verdoppelt werden, bereits verinnerlicht haben. Verwechslungen zwischen <sch> und <ch> im Silbengelenk sind einerseits auf phonologische Unterscheidungsschwierigkeiten zurückzuführen (vgl. Kapitel 4.2.5). Andererseits belegt die ebenenübergreifende Analyse mit der Annotationsebene [Morph_AffixS], dass ein beträchtlicher Anteil dieser Verwechslungen (61,90 %; n = 13) mit der korrekten Schreibung der Suffixe {-isch} und {-lich} zusammenhängt, und somit morphologisch begründet werden kann.

Aus dieser Analyse ergibt sich, dass die Abweichungen am Endrand, die chinesischen Deutschlernenden erhebliche Schwierigkeiten bereiten, mit der komplexen Phonetik und den orthographischen Konventionen des deutschen Silbenauslauts zusammenhängen. Starke Transfereffekte aus L1 (insbesondere die Tendenz zur offenen Silbe oder zum einfachen Nasalauslaut) sowie die Herausforderungen durch spezifisch deutsche Phänomene wie Auslautverhärtung und Silbengelenkschreibung prägen das Fehlerbild.

4.4 Der morphologische Bereich

Die morphologische Schreibung im Deutschen basiert auf dem Prinzip der morphologischen Konstanz. Dieses Prinzip besagt, dass Wortstämme in verschiedenen Formen, Ableitungen und Komposita eine möglichst einheitliche oder ähnliche graphemische Gestalt bewahren sollen, auch wenn sich die lautliche Realisierung ändern kann. Als ein leserfreundliches Schriftsystem (vgl. Eisenberg, 2017) strebt die deutsche Orthographie danach, die morphologische Verwandtschaft von Wörtern auch graphemisch sichtbar zu machen. Ein klassisches Beispiel hierfür ist das Wortpaar <apfel> – <äpfel>. Obwohl der Umlaut /ɛ/ in <Äpfel> phonetisch auch mit <e> verschriftet werden könnte, wird hier <ä> verwendet, um die morphologische Beziehung zum Singular <Apfel> (mit <a>) zu verdeutlichen (vgl. Fuhrhop, 2021, S. 25).

Die Bedeutung der morphologischen Schreibung für die orthographische Kompetenz spiegelt sich auch in etablierten Rechtschreibdiagnoseverfahren wider. In der HSP werden beispielsweise Rechtschreibphänomene wie Auslautverhärtung, abzuleitende Umlautungen, Kompositumschreibung und Präfixschreibung als Teil der morphologischen Strategie betrachtet (vgl. May, 2013, S. 16). Die AFRA differenziert Fehlerkategorien auf der morphologischen Ebene wie morphologische Segmentierung, Morphem-Differenzierung, Fehler bei unselbständigen Morphemen sowie bei konsonantischer als auch bei vokalischer Ableitung (vgl. Herné & Naumann, 2016, S. 12–14). In der OLFA wird die morphologische Schreibung zwar den Fehlerkategorien in der orthographischen Phase (Gruppe III) zugeordnet (vgl. Thomé & Thomé, 2020, S. 20–26), aber es findet keine spezifische, explizit auf dem morphologischen Prinzip basierende Fehlerkategorisierung statt.

In der vorliegenden Untersuchung werden die orthographischen Abweichungen im morphologischen Bereich mithilfe folgender Annotationsebenen analysiert:

- [Morph_Fehler] und [Morph_ZH]: die grundlegende Segmentierung des Lerner-texts bzw. der Zielhypothese nach Morphemen.
- [Morphem]: Abweichungen zwischen [Morph_Fehler] und [Morph_ZH] werden nach der Art des betroffenen Morphems klassifiziert (lexikalisches Morphem (LM) vs. grammatisches Morphem (GM)).
- [Morph_Konstanz]: Fehler bei der Stammkonstanzschreibung (z. B. bei Umlautung, Auslautverhärtung usw.).
- [Morph_KompoS]: Fehler bei der Kompositumschreibung (z. B. Fugenelemente usw.).
- [Morph_AffixS]: Fehler bei der Affixschreibung.

Es ist anzumerken, dass Abweichungen, die nicht spezifische Aspekte der morphologischen Schreibung wie Stammkonstanz, Kompositions- oder Affixschreibung betreffen (diese werden in den folgenden Unterkapiteln 4.3.1 und 4.3.2 detailliert analysiert), auf der Annotationsebene **[Morphem]** lediglich generalisierter als Fehler bei lexikalischen Morphemen (LM) oder grammatischen Morphemen (GM) erfasst werden. Lexikalische Morpheme tragen hierbei die Hauptbedeutung eines Wortes und sind in der Regel eigenständige Bestandteile des Wortschatzes. Grammatische Morpheme hingegen dienen der Darstellung grammatischer Beziehungen und Funktionen. Sie tragen keine eigenständige lexikalische Bedeutung, sondern drücken Informationen wie Kasus, Numerus, Genus, Tempus, Modus aus. In dieser Annotationsebene [Morphem] werden die Abweichungen je nach der Fehlerart als Hinzufügung, Auslassung oder Ersetzung zusätzlich zu den betroffenen Morphemarten markiert (z. B. eine Auslassung bei einem grammatischen Morphem wird als GM-annotiert).

Im gesamten DeChiLKo-Korpus wurden auf der Ebene **[Morphem]** 3.616 fehlerhafte Morpheme in 327 Lernertexten identifiziert. Die Verteilung dieser Fehler nach Morphemarten (lexikalische vs. grammatische Morpheme) und Fehlerarten (Auslassung, Hinzufügung, Ersetzung) wird in Abbildung 26 dargestellt.

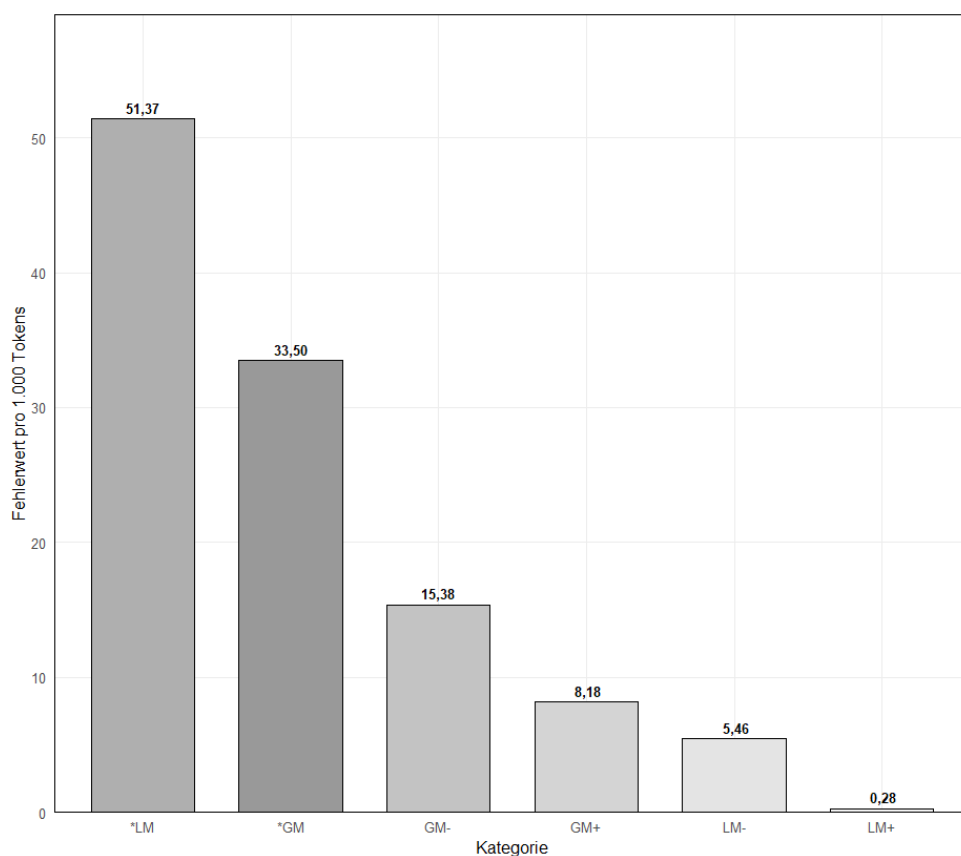


Abbildung 26: Fehlerverteilungen nach Morphemarten; Fehlerwert pro 1.000 Tokens

Bemerkenswert ist, dass die Fehlerhäufigkeit bei lexikalischen Morphemen insgesamt ähnlich hoch ausfällt wie bei grammatischen Morphemen. Dies deutet darauf hin, dass die Lernenden sowohl bei der Wortstammschreibung als auch bei den grammatischen Markierungen Schwächen aufweisen. Hervorzuheben ist jedoch, dass falsche lexikalische Morpheme (*LM) mit einem Fehlerwert von 51,37 pro 1.000 Tokens am häufigsten im Korpus vorkommen. Anschließend folgen fehlerhafte grammatische Morpheme (*GM) mit einem Fehlerwert von 33,50 pro 1.000 Tokens. An dritter Stelle stehen fehlende grammatische Markierungen (GM-) mit 15,38 pro 1.000 Tokens. Weniger häufig tritt die überflüssige Hinzufügung grammatischer Morpheme (GM+) auf. Die Fehlerkategorie „überflüssige lexikalische Morpheme“ (LM+) kommt hingegen im Korpus nur sehr selten vor, was durch den Aufgabentyp – ein Diktat – plausibel erklärbar ist.

Die spezifischen Abweichungen, die die morphologische Schreibung im engeren Sinne betreffen, werden in den folgenden Unterkapiteln detailliert auf Basis der Anno-

tationsebenen für die Stammkonstanzschreibung ([Morph_Konstanz], Kapitel 4.4.1) sowie die Wortbildung ([Morph_KompoS] und [Morph_AffixS], Kapitel 4.4.2) analysiert.

4.4.1 Fehler bei der morphologischen Stammkonstanz [Morph_Konstanz]

Auf der Annotationsebene [Morph_Konstanz] wird die spezifische morphologische Schreibung im Hinblick auf die Einhaltung der Stammkonstanz geprüft. Die dabei identifizierten Abweichungen werden in fünf Hauptfehlerkategorien unterteilt. Erstens werden Abweichungen bei der morphologisch bedingten Umlautung (UL) betrachtet, wie beispielsweise *<Epfel> statt *Äpfel* bei der Pluralform. Zweitens umfasst die Kategorie Auslautverhärtung (ALV) Fehler bei der Schreibung von auslautenden Konsonanten, wie *<Hunt> statt *Hund* (/hont/). Darüber hinaus werden Abweichungen bei übernommenen silbischen Schreibungen unterschieden. Dazu zählen:

- Fehler bei morphologischen Konsonantenverdopplungen (KD), wie z. B. *<Altag> statt *Alltag*;
- Fehler bei Dehnungsschreibungen (DH), bei denen ein Dehnungs-<h> zur Wahrung der Stammkonstanz auch in flektierten oder abgeleiteten Formen erhalten bleibt (z. B. *<fült> statt *fühlt*);
- Fehler bei der Schreibung des silbeninitialen <h> (SinH), wenn dieses morphologisch relevant ist (z. B. *<drest> statt *drehst*).

Darüber hinaus fasst die Kategorie *Stamm* alle anderen Verstöße gegen die Stammkonstanzschreibung zusammen, die sich nicht eindeutig einer der oben genannten Kategorien zuordnen lassen, jedoch die graphemische Konstanz des Wortstammes betreffen.

Insgesamt wurden 292 Abweichungen hinsichtlich der Stammkonstanzschreibung in 166 Lernertexten im DeChiLKO-Korpus gefunden. Dies entspricht einem Fehlerwert von 9,22 pro 1.000 Tokens (vgl. Tabelle 33).

Tabelle 33: Überblick über Abweichungen der Stammkonstanzschreibung

Abweichungen der Stammkonstanzschreibung insgesamt	Summe aller Tokens	Abweichungen der Stammkonstanzschreibung pro 1.000 Tokens
292	31.674	9,22

Die Fehlerverteilung innerhalb dieser Kategorie (vgl. Abbildung 27) legt dar, dass die morphologische Umlautung (UL) mit 2,37 Fehlern pro 1.000 Tokens die häufigste Fehlerquelle darstellt. Es ist zu beachten, dass Mehrfachkodierungen bei der Annotation möglich waren, wenn ein Fehler mehrere Aspekte der Stammkonstanz berührte. Beispielsweise wurde die Fehlschreibung des Morphems {*mäß*} als <mess> statt <*mäß*> im Wort *regelmäßig* sowohl als Fehler bei der morphologisch ableitenden Umlautung (Derivation vom *Basisstamm* {*maß*} zu {-*mäßig*}) als auch als Verstoß gegen die *ß*-Konstanz dieses Stammes {*maß*} gewertet und daher als „UL/Stamm“ markiert.

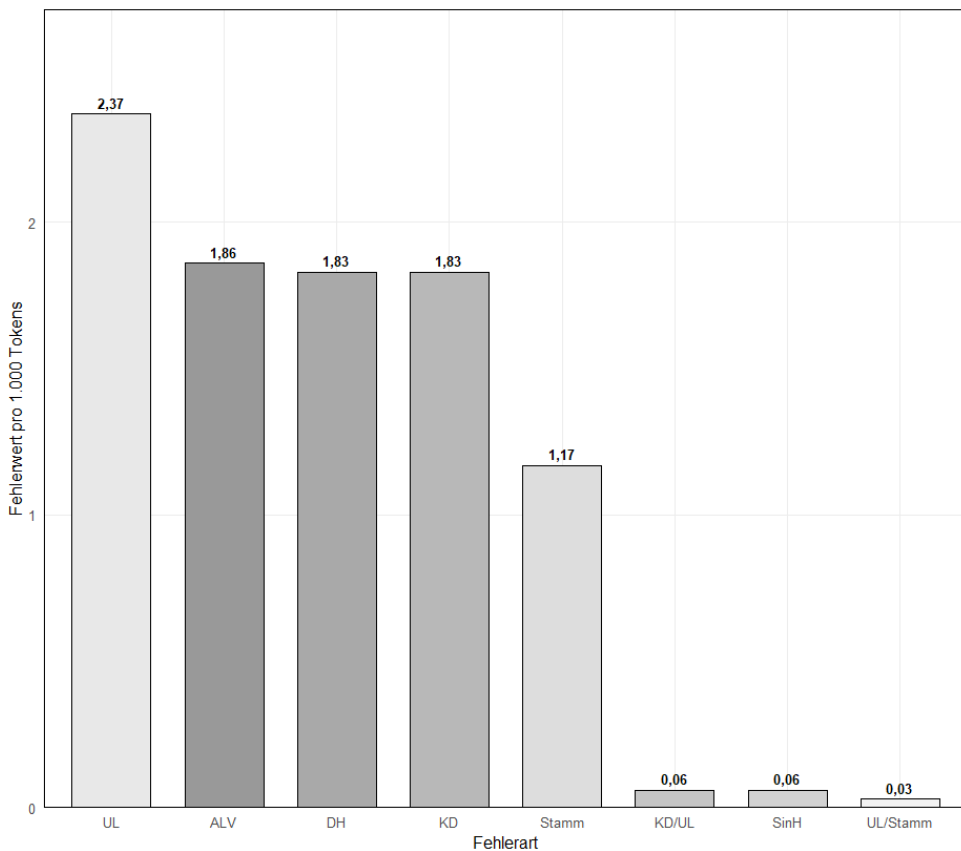


Abbildung 27: Fehlerverteilung bei der Stammkonstanzschreibung; Fehlerwert pro 1.000 Tokens

Bemerkenswert ist die hohe Fehleranfälligkeit bei der morphologisch bedingten Umlautung (UL). Im Deutschen tritt diese in vielfältigen Kontexten der Flexionsmorphologie auf (z. B. Pluralbildung von Substantiven: *Väter, Äpfel, Gründe, Länder*; Adjektivkomparation: *älter, ärmer*; Verbflexion: *fährt, hätte*), sowie in der Derivationsmorphologie (z. B. *backen – Bäcker*) (vgl. Fuhrhop & Peters, 2013, 240 f.; Fuhrhop, 2021). Die ebenenübergreifende Analyse mit der Part-of-Speech-Annotation der Zielhypothese ([ZHpos]) im DeChiLKo deutet darauf hin, dass Abweichungen bei der morphologischen Umlautung vor allem bei Substantiven ($n = 42$) und Verben ($n = 24$) zu finden sind, während Abweichungen bei Adjektiven seltener sind ($n = 8$). Eine nähere Betrachtung dieser Abweichungen macht deutlich, dass die Lernenden insbesondere Schwierigkeiten mit der Umlautung im Plural haben ($n = 27$); die meisten Fehler treten hier bei der Pluralbildung von *Grund* auf (z. B. **Grund* statt *Gründe*) auf.

Auch die Umlautung bei der Verbflexion ($n = 20$) stellt eine Lernhürde dar. Hier weisen die Lernenden vor allem Probleme bei den Modalverben nach, was insbesondere für *müssen* gilt: Statt *muss* wird häufig **müss* für die Singularform und **mussen* als Infinitivform geschrieben. Diese Fehlschreibungen lassen sich deswegen als Überge-

neralisierungsfehler interpretieren. In der deutschen Verbflexion scheint es häufig eine „Richtung“ der Umlautung vom Infinitiv zur 3. Person Singular zu geben (z. B. *fahren* – *fährt*, *laufen* – *läuft*, *raten* – *rät*), und Präsensformen sind oft nicht komplexer als Präteritumformen (vgl. Eisenberg, 2016, S. 80). Es ist denkbar, dass die Lernenden diese beobachtete Regularität übergeneralisieren und fälschlicherweise den Infinitiv von *müssen* als **mussen* und die 3. Person Singular als **müss* schreiben. Eine ähnliche Übergeneralisierung findet sich auch bei der Genitivschreibung von *Land*: Statt *des Landes* wird in den Lernertexten **des Länder* oder **des Ländes* geschrieben. Da bei der Pluralbildung von *Land* eine Umlautung stattfindet (*die Länder*), könnten die Lernenden diese erworbene Regel fälschlicherweise auch auf die Genitivbildung übertragen haben. Im Vergleich zur Pluralmarkierung und Verbflexion sind Umlautungen bei Derivationen und in der Adjektivkomparation im DeChiLKo weniger fehlerauffällig.

Diese Fehlerphänomene bei der Umlautung sind im Kontext des multikompetenten Repertoires chinesischer Lernenden besonders relevant. Die L1 Chinesisch charakterisiert sich als eine isolierende Sprache ohne Flexionsmorphologie und ohne Konzept einer Umlautung zur Markierung grammatischer Funktionen. Auch die erste Fremdsprache Englisch kennt die deutsche Form der Umlautung nicht. Die Lernenden müssen dieses zielsprache-spezifische System gänzlich neu erwerben, was die hohe Fehleranfälligkeit und die Tendenz zur Übergeneralisierung plausibel erklärt.

An zweiter Stelle der Fehlerhäufigkeit treten Abweichungen bei der Auslautverhärtung (ALV) mit 1,86 Fehlern pro 1.000 Tokens hervor. Im Deutschen können im Silbenendrand phonetisch keine stimmhaften Obstruenten stehen; orthographisch wird jedoch im Sinne der Morphemkonstanz in allen Formen das Graphem für den zugrundeliegenden stimmhaften Obstruenten beibehalten (vgl. Eisenberg, 2016, S. 83). Die korrekte Schreibung trotz Auslautverhärtung gilt als fehleranfällig und als Indikator für Förderbedarf im Orthographieerwerb (vgl. Corvacho Del Toro, 2016, S. 403; Nimz, 2022, S. 39). Im DeChiLKo-Korpus fallen Abweichungen bei der Auslautverhärtung vor allem in den Morphemen **<bat>* für *<bad>* und **<rat>* für *<rad>* auf. Auch die Ersetzung von *<g>* durch *<k>* ist häufig (z. B. **<tak>* für *<tag>*, **<lik>* für *<leg>*) belegt. Bemerkenswert ist, dass das Graphem *<g>* am Endrand oft als *<ck>* verschriftet wurde (z. B. **<weck>* für *<weg>*, **<leck>* für *<legt>*, **<krick>* für *<krieg>*). Solche Fehlschreibungen deuten einerseits darauf hin, dass die Lernenden bereits mit spezifischen Phonem-Graphem-Korrespondenzen des Deutschen vertraut sind, die nicht zu den Basisgraphemen zählen (vgl. Thomé et al., 2016). Andererseits legen diese Fehlschreibungen nahe, dass eine fehlerhafte Perzeption der Vokalquantität (kurzer Vokal vor *<ck>*) einen negativen Einfluss auf die Schreibung des gesamten Morphems ausüben kann. Übergeneralisierungsfehler bei der Auslautverhärtung (z. B. **<Prozend>* statt *Prozent*) sind im Korpus hingegen selten.

Die übernommenen silbischen Schreibungen, insbesondere Fehler beim Dehnungs-*<h>* (DH), stellen ebenfalls eine wichtige Fehlerquelle für die chinesischen Deutschlernenden dar (Fehlerwert 1,83/1.000 Tokens). Im Sinne der morphologischen Konstanz bleibt das Dehnungs-*<h>* auch in solchen Formen erhalten, in denen dem betonten Vokal mehrere Konsonanten folgen (vgl. Eisenberg, 2016, S. 81). Im Deutschen

ist die Schreibung mit Dehnungs-*<h>* besonders bei den Verben wichtig, um auch in flektierten Formen mit komplexem Endrand die Länge bzw. Gespanntheit des Vokals zu signalisieren (vgl. Fuhrhop, 2021, S. 27). Die Analyse der POS-Tags der Zielhypothese [ZHpos] zeigt, dass Auslassungen des morphologisch relevanten Dehnungs-*<h>* im Korpus auch sehr häufig bei Substantiven ($n = 53$) vorkommen. Am häufigsten wird das Morphem {lohn} im Wort *Belohnung* als *<lon>* geschrieben; auch die Fehlschreibung **<far>* für das Morphem {fahr} im Wort *Fahrrad* tritt häufig auf. In diesen beiden Wörtern, *Belohnung* und *Fahrrad*, ist das *<h>* aus rein phonographischer Sicht für die Markierung der Vokallänge gewissermaßen überflüssig, da es sich um eine betonte Silbe mit einfachem Endrand handelt und der Vokal daher auch ohne *<h>* lang gelesen wird. Die beiden Substantive werden jedoch trotzdem mit *<h>* geschrieben, da sie Derivate bzw. Komposita sind und die morphologische Konstanz zum Stamm ({lohn}, {fahr}) gewahrt werden soll. In den Verbformen *lohn*en oder *fahr*en ist das *<h>* nicht primär für die Vokallänge notwendig; jedoch enthalten die Verbformen das Dehnungs-*<h>*. In flektierten Formen wie *lohn*st oder *fähr*st trägt das *<h>* zur Kennzeichnung der Vokallänge bei. Gerade diese doppelte Komplexität – die oft nicht direkt aus der Lautung ableitbare und morphologisch bedingte Setzung des Dehnungs-*<h>* – stellt eine Lernhürde für DaF-Lernende dar, da in ihren bestehenden multiplen Repertoires kein vergleichbares Graphem bekannt ist.

Abweichungen bei der morphologischen Konsonantenverdopplung (KD) treten ebenfalls häufig auf (Fehlerwert: 1,83 /1.000 Tokens). Die Doppelkonsonantschreibung im Deutschen ist primär phonologisch begründet und dient als Silbengelenk zwischen einem betonten Vokal und einem unbetonten Vokal (vgl. Eisenberg, 2016, S. 47). Daher ist die Doppelkonsonanz in Einsilbern im heutigen Deutsch ausschließlich morphologisch bedingt (vgl. Fuhrhop & Peters, 2013, S. 242); die Gelenkschreibung bleibt also erhalten, auch wenn der entsprechende Konsonant in der einsilbigen Form nicht mehr als Silbengelenk fungiert (vgl. Eisenberg, 2016, S. 102). Im DeChiLKo-Korpus treten die meisten Abweichungen in dieser Kategorie bei Substantiven ($n = 42$) auf, wobei Fehler besonders bei Doppelkonsonanten in Komposita ($n = 33$) auffällig sind. Häufig wird **<Altag>* statt *<Alltag>* geschrieben. Auch Fehlschreibungen wie **<schwim>* für *<schwimm>* im Kompositum *Schwimmbad* sind verbreitet. Des Weiteren wiesen die Lernenden Schwierigkeiten mit der morphologischen Gelenkschreibung in der Verbflexion auf (z. B. **<solte>* für *<sollte>*; **<geschmect>* für *<geschmeckt>*; **<jogt>* für *<joggt>*). Bemerkenswert ist, dass die Auslassung der Doppelkonsonanz häufig zusammen mit der Verwechslung des Silbenkerns durch einen Diphthong auftritt (z. B. **<Eiszimmer>* für *<Esszimmer>*, **<streiß>* für *<Stress>*, **<Schweim>* für *<Schwimm>*). Solche Fehlschreibungen deuten wiederholt darauf hin, dass die chinesischen Lernenden die Grundregel der Markierung von Vokalgespanntheit zwar bereits verinnerlicht zu haben scheinen, das Prinzip der Stammkonstanz aber noch nicht sicher beherrschen.

Weniger häufig sind Fehler, die andere Aspekte der Stammkonstanzschreibung (Stamm) betreffen. Hierunter fallen im DeChiLKo alle Abweichungen bei Adjektiven, wie z. B. **<tächnisch>* statt *<technisch>*. Obwohl für das Phonem /ɛ/ sowohl das Gra-

phem <ä> als auch <e> möglich ist, setzt die Stammkonstanz hier die Beibehaltung des Stamms {techn} auch in der Derivation voraus. Abweichungen beim morphologischen silbeninitialen <h> (SinH), wie etwa *<stet> für <steht>, sind im Korpus sehr selten zu finden. Ebenfalls selten treten Abweichungen auf, die auf mehrere Fehlerquellen der morphologischen Schreibung gleichzeitig zurückzuführen sind.

Interessanterweise wird Instabilität bei der Verschriftlichung desselben Wortstammes im DeChiLKo an mehreren Stellen beobachtet. Beispielsweise wird das Dehnungs-<h> im Verb *fahren* von einigen Lernenden korrekt realisiert, beim Kompositum *Fahrrad* jedoch von denselben Lernenden ausgelassen und als *<Farrad> geschrieben. Solche instabilen Verschriftungen deuten darauf hin, dass ein gefestigtes lexikalisch-semanticches Bewusstsein für die Wortidentität über verschiedene morphologische Kontexte hinweg noch nicht vollständig aufgebaut ist (vgl. Hans-Bianchi & Katelhön, 2010, S. 46). Dieses Schwanken zwischen korrekten und inkorrekten Formen desselben Stammes unterstreicht zudem die Nicht-Linearität der Rechtschreibentwicklung.

Zusammenfassend lässt sich für den Bereich der Stammkonstanz festhalten, dass die chinesischen Deutschlernenden erhebliche Schwierigkeiten mit den morphologisch bedingten Schreibvarianten im Deutschen haben. Dies betrifft insbesondere die Umlautung, die Auslautverhärtung sowie die Beibehaltung von Dehnungs-<h> und Konsonantenverdopplungen im Stamm. Die beobachteten Übergeneralisierungen deuten auf aktive, aber noch unvollständige Regelbildungsprozesse hin. Die Beobachtung der Instabilitäten legt nahe, dass der Aufbau eines konsistenten orthographischen Wissens, insbesondere bei morphologisch komplexen Wörtern, ein dynamischer und von Variabilität geprägter Prozess ist.

4.4.2 Fehler bei der Wortbildung: [Morph_Kompo] [Morph_Affix]

Die deutsche Wortbildung kennt vier Grundtypen: Komposition, Präfigierung, Suffigierung und Konversion. Die Komposition verbindet zwei oder mehrere Wortstämme, die einander nebengeordnet sind. Bei der Präfigierung wird ein Wortbildungsaffix vor den Stamm gesetzt, bei der Suffigierung wird der Wortstamm durch ein nachgestelltes Suffix ergänzt. Die Konversion, bei der ein Wort in eine andere Wortart überführt wird, ermöglicht die Verwendung eines Wortbildungsaffixes vor den Stamm. Bei der Suffigierung wird dem Wortstamm hingegen ein nachgestelltes Suffix hinzugefügt (vgl. Eisenberg, 2013, 201 f.). In der vorliegenden Arbeit liegt der Fokus auf der Komposition, Präfigierung und der Suffigierung. Abweichungen in diesen Wortbildungsprozessen werden auf den Annotationsebenen [Morph_Kompo] und [Morph_Affix] erfasst. Abweichungen, die primär auf Konversionsprozesse zurückzuführen sind und sich oft in der Groß- und Kleinschreibung manifestieren (z. B. die Substantivierung von Verben), werden im Rahmen der syntaktischen Fehleranalyse in Kapitel 4.5.1 genauer betrachtet.

Die Fugenbildung gilt als eines der wichtigsten morphologischen Mittel zur Strukturierung von Komposita im Deutschen (vgl. Fuhrhop, 2000, S. 212). Auf der Annotationsebene [Morph_Kompo] konzentriert sich die Analyse vor allem auf Abweichungen bei diesen Fugenelementen. Sie werden als sprachliche Einheiten definiert, durch die sich das Erstglied eines Kompositums von seiner isolierten Grundform un-

terscheidet – was bei substantivischen Erstgliedern in der Regel einer Abweichung von der Nominativ-Singular-Form entspricht. Als typische Fugenelemente im Deutschen gelten {-s-}, {-n-}, {-en-}, {-er-} und {-e-} (vgl. Fuhrhop, 2000, S. 202). Fehler bei Fugenelementen werden nach ihrer Art als überflüssiges Fugenelement (Fugen+), ausgelassenes Fugenelement (Fugen-) oder falsches Fugenelement (*Fugen) kategorisiert.

Ein weiteres relevantes Phänomen bei der Kompositaschreibung ist die Tilgung von Lauten an der Morphemgrenze, insbesondere die Reduktion identischer Konsonanten an der Morphemfuge („Konsonantenreduktion“ oder „Phonemverschmelzung“). Beispielsweise wird *annehmen* (/ʔan,ne:mən/) in der Umgangslautung als /ʔane:mən/ realisiert (Eisenberg, 2016, S. 79). Im Geschriebenen kann eine solche Phonemverschmelzung zur fehlerhaften Auslassung eines Konsonantengraphems führen (z. B. *<anehmen>). Im DeChiLKo-Korpus werden solche Abweichungen als Phonemverschmelzung (PV) annotiert.

Insgesamt wurden 134 Abweichungen bei der Kompositumschreibung im DeChiLKo-Korpus gefunden, was einem Fehlerwert von 4,23 pro 1.000 Tokens entspricht (vgl. Tabelle 34).

Tabelle 34: Übersicht über Abweichungen bei der Kompositaschreibung

Abweichungen Kompositumschreibung insgesamt	Summe aller Tokens	Abweichungen bei der Wortschreibung pro 1.000 Tokens
134	31.674	4,23

Die folgende Abbildung 28²⁴ informiert über die genaue Fehlerverteilung bei der Kompositaschreibung je nach Fehlerart.

24 Im Gegensatz zur allgemeinen Tokennormierung bezieht sich der Fehlerwert in der Kategorie Kompositaschreibung auf die Gesamtzahl der produzierten Komposita (n = 1.670). Dies isoliert die spezifische Fehleranfälligkeit der Wortbildung von der bloßen Vorkommenshäufigkeit der Struktur im Korpus.

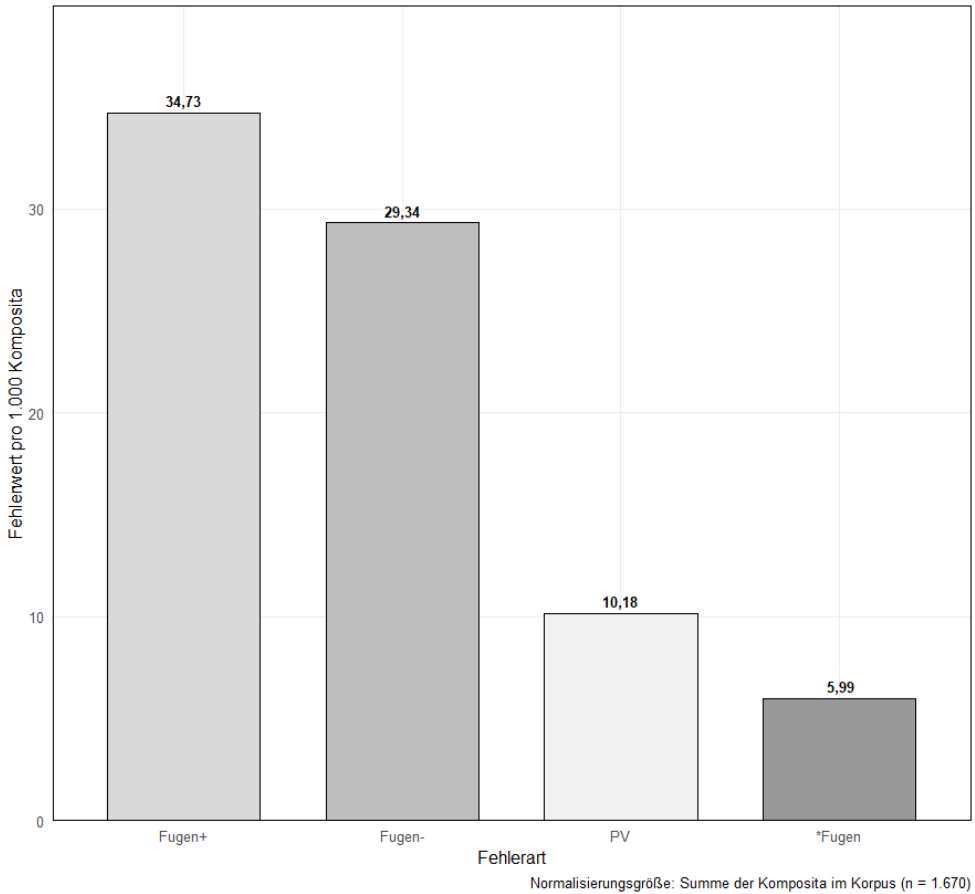


Abbildung 28: Fehlerverteilung nach Fehlerarten bei der Kompositaschreibung; Fehlerwert pro 1.000 Komposita; Normalisierungsgröße: Summe der Komposita im Korpus (n = 1670)

Die Fehlerverteilung belegt, dass Abweichungen im Umgang mit Fugenelementen den größten Anteil ausmachen, während Fehler durch Phonemverschmelzung (PV) seltener auftreten (Fehlerwert 10,18/1.000 Komposita). Innerhalb der Fehlerkategorien bei Fugenelementen sind überflüssige Fugenelemente (Fugen+) mit einem Fehlerwert von 34,73 pro 1.000 Komposita die häufigste Fehlerart. Dies bedeutet, dass die chinesischen Lernenden in einigen Fällen zusätzliche Elemente zwischen zwei Stämmen einfügen, obwohl diese orthographisch nicht erforderlich sind. Am häufigsten ist das Fugenelement {-en-} hinzugefügt (n = 27) festgestellt worden, wie z. B. in *<Schwimmenbad> statt <Schwimmbad> (n = 20) oder *<Lernendauer> statt <Lerndauer> (n = 4). Auch die Fugen {-e-} (n = 13) und {-er-} (n = 11) werden häufig als überflüssiges Element in Komposita eingesetzt (z. B. *<Lehrerinhalte> statt <Lehrinhalte> (n = 8); *<Spieledauer> statt <Spieldauer>). Obwohl die Mehrzahl der Komposita im deutschen Wortschatz fugenlos gebildet wird: 73 % der Substantiv- und 70 % der Adjektivkomposita haben eine

Nullfuge(vgl. Ortner et al., 1991, S. 54), kann die Komplexität der Kompositabildung mit unterschiedlichen Fugen dazu führen, dass die Notwendigkeit eines Fugenelements übergeneralisiert und fälschlicherweise eines eingefügt wird.

Fehlende Fugenelemente (Fugen-) treten mit einer Häufigkeit von 29,34 pro 1.000 Komposita auf und stellen die zweithäufigste Fehlerart dar. Die Lernenden scheinen hierbei Schwierigkeiten zu haben, die Notwendigkeit eines Fugenelements zu erkennen und dieses korrekt umzusetzen. Am häufigsten geschehen die Auslassungen des Fugen-{s} (n = 20), wie z. B. in *<Reichstaggebäude> und *<Geburtstagspielen>, sowie des Fugenelements {-en-} in Komposita wie z. B. *<Wochende> für <Wochenende> (n = 18). Bei den Fehlschreibungen *<Geburtstagsspiele> und *<Wochende> handelt es sich um Fälle, in denen durch das Fugenelement eine Duplikation des Graphems <s> bzw. der Graphemfolge <en> entstünde (*Geburtstag-s-spielen*, *Woch-en-ende*). Es ist durchaus denkbar, dass die Lernenden hier aus Vereinfachung das Fugenelement wegließen.

Die häufigen Weglassungen und Hinzufügungen von Fugenelementen könnten auf den Einfluss der Kompositabildung in der L1 Chinesisch zurückzuführen sein. Die chinesische Sprache verfügt bei Substantivkomposita über ein eigenes, jedoch mit dem Deutschen vergleichbares Wortbildungsverfahren (vgl. Ma, 1984, S. 52). Die Komposita im Chinesischen lassen sich in fünf strukturelle Hauptkategorien unterteilen (vgl. Karlgren, 2001, S. 23–25):

1. koordinative Komposita, bei denen die Bestandteile gleichrangig sind, wie z. B. 父母 (*fùmǔ*, „Eltern“, wörtlich „Vater und Mutter“)
2. Determinativkomposita, bei denen ein Bestandteil den anderen näher bestimmt, wie z. B. 大山 (*dàshān*, „großer Berg“, wörtlich „groß + Berg“)
3. komplementäre Komposita, bei denen ein Bestandteil den anderen ergänzt, wie etwa 学习 (*xuéxí*, „lernen“, wörtlich „studieren + üben“)
4. Verb-Objekt-Komposita, bei denen ein Verb mit einem Objekt kombiniert wird, z. B. 看书 (*kànshū*, „lesen“, wörtlich „lesen + Buch“)
5. Subjekt-Prädikat-Komposita, bei denen ein Bestandteil als Subjekt und der andere als Prädikat fungiert, z. B. 心痛 (*xīntòng*, „das Herz schmerzt“, wörtlich „Herz + Schmerz“)

Wie im Deutschen stellt die Substantivkomposition auch im Chinesischen ein äußerst produktives Wortbildungsverfahren dar. Besonders häufig treten die Substantivkomposita in chinesischen juristischen Texten auf (vgl. Zhu & Best, 1992, S. 56). Entscheidend ist jedoch, dass Chinesisch eine isolierende Sprache ohne Flexion ist. Das bedeutet, dass die Bestandteile eines Kompositums meist unverändert und ohne explizite Fugenelemente nebeneinandergestellt werden, wie z. B. beim deutschen Kompositum *Wochenende*, das im Chinesischen als 周末 (*zhōumò*, wörtlich: Woche + Ende) realisiert wird. Die Notwendigkeit, im Deutschen spezifische, oft unregelmäßige Fugenelemente einzusetzen, mag daher eine erhebliche Lernhürde für chinesische Deutschlernende darstellen.

Neben diesem L1-Einfluss trägt auch die inhärente Unregelmäßigkeit der Verwendung von Fugenelementen im L3-Deutsch zur Fehleranfälligkeit bei. Obwohl Fugenelemente oft als Flexionsmarker und insbesondere als solche von Substantiven vorkommen, weichen die Fugenformen häufig von den regulären Flexionsformen ab (vgl. Eisenberg, 2016, S. 226). Beispielsweise sind subtraktive Fugen (wie in *Geschichtsbuch*) keine typischen Flexionssuffixe. Zudem treten sogenannte „semantisch falsche Plurale“ (*Gänsebraten*) oder „semantisch falsche Singulare“ (*Freundeskreis*) auf. Auch die fehlende *s/es*-Variation bei der *s*-Fuge im Vergleich zur Genitivmarkierung (z. B. <Jahresende> vs. *<Jahresende>; <Anwaltskammer> vs. *<Anwalteskammer>) ist für Lernende schwer nachvollziehbar (vgl. Wegner, 2003, S. 431). Solche Inkonsistenzen im Zielsprachsystem führen dazu, dass Lernende die Verwendung der Fuge oft als unvorhersehbar empfinden und Schwierigkeiten haben, das korrekte Fugenelement zu identifizieren und einzusetzen.

Die Auslassung bei der Phonemverschmelzung (PV) (10,18/1.000 Komposita) tritt vor allem beim Kompositum *Fahrrad* auf, wo die an der Morphemgrenze aufeinander treffenden <rr> zu einem einfachen <r> reduziert werden. Im Gegensatz zu *annehmen*, wo die Reduktion in der Umgangssprache üblich ist, findet bei Kompositions-fugen im Allgemeinen keine Reduktion statt (z. B. *Türrahmen*, *Lauffeuer*), auch nicht bei drei gleichen Konsonanten (*Werkstatttreppe*) (vgl. Eisenberg, 2016, S. 79). Konsonantenreduktionen sind im Deutschen an der Grenze zwischen Stamm und Flexionsendung grammatikalisiert (*reisen* – *du reist* statt **reisst*). Diese unterschiedliche Handhabung der Konsonantenreduktion je nach morphologischem Kontext kann für DaF-Lernende ein zusätzliches Lernhindernis darstellen.

Die Verwechslung von Fugenelementen (*Fugen) kommt vergleichsweise selten vor. Die Analyse der falsch eingesetzten Fugenelemente verdeutlicht, dass die Lernenden vor allem das Fugenelement {-e-} mit {-er-} verwechseln, z. B. *<Bundersbürger> für *Bundesbürger*. Der Fehler *<Bundersbürger> könnte darauf hindeuten, dass der Lernende das Prinzip der Fugenelementsetzung in deutschen Komposita teilweise bereits verstanden hat, also eine Notwendigkeit für ein Verbindungselement zwischen den Kompositionsgliedern erkennt. Jedoch könnte es zutreffen, dass aufgrund des auditiven Inputs im Rahmen der Diktataufgabe eine fehlerhafte Segmentierung oder Identifizierung ähnlich klingender Lautelemente vorgenommen wurde, die dann zu einer inkorrekten graphemischen Realisierung des Fugenelements {-es-} als vermeintliches {-ers-} führte.

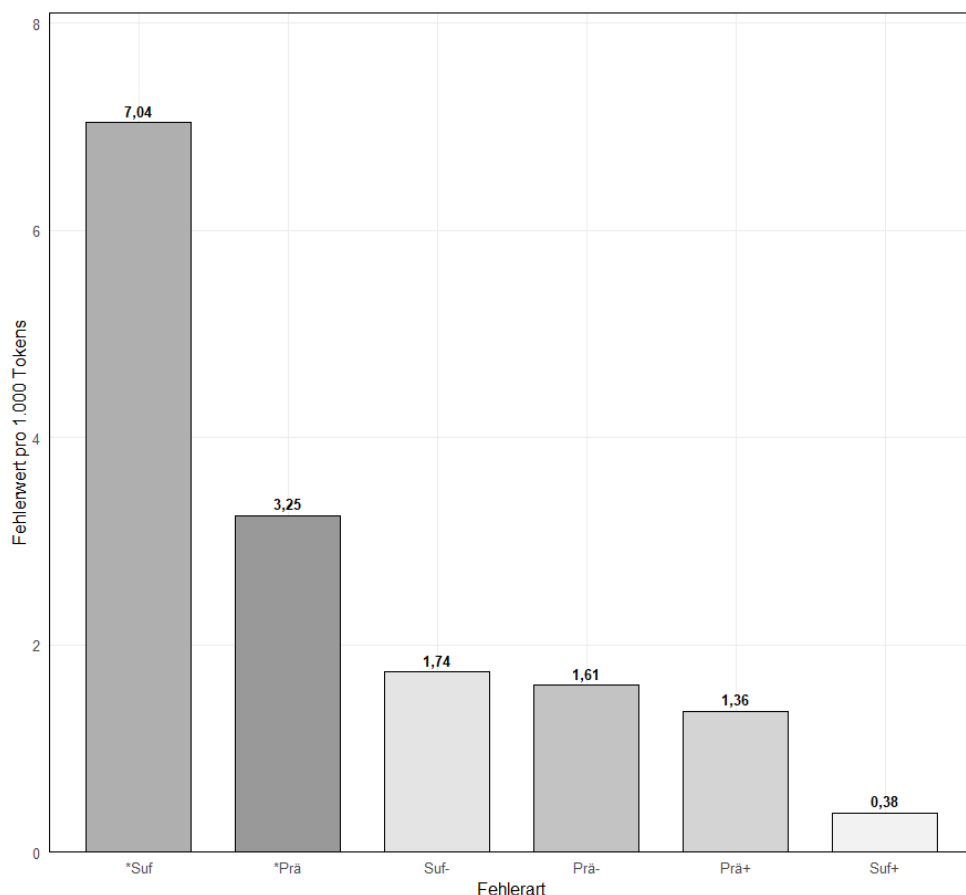
Auf der Annotationsebene [Morph_Affix] stehen die Abweichungen bei der Affixschreibung im Fokus, wobei primär Derivationsaffixe betrachtet werden. Es wurde zwischen Präfixen (Prä) und Suffixen (Suf) sowie nach den Fehlerarten Auslassung (-), Hinzufügung (+) und Verwechslung (*) unterschieden.

Insgesamt wurden 487 Abweichungen hinsichtlich der Affixschreibung im Korpus gefunden, was einem Fehlerwert von 15,38 pro 1.000 Tokens entspricht (vgl. Tabelle 35). Im Vergleich zur Kompositaschreibung (Fehlerwert 4,23) weisen die Lernenden somit deutlich mehr Schwierigkeiten bei der Affixschreibung auf.

Tabelle 35: Übersicht über Abweichungen bei der Affixschreibung

Abweichungen bei der Affixschreibung insgesamt	Summe aller Tokens	Abweichungen bei der Wortschreibung pro 1.000 Tokens
487	31.674	15,38

Die Fehlerverteilung bei der Affixschreibung (vgl. Abbildung 29) macht deutlich, dass Abweichungen bei Suffixen häufiger als bei Präfixen auftreten. Die häufigste Fehlerart ergab falsche Suffixe (*Suf) mit einem Fehlerwert von 7,04 pro 1.000 Tokens. An zweiter Stelle stehen falsche Präfixe (*Prä) mit einem Fehlerwert von 3,24 pro 1.000 Tokens. Weniger häufig sind Auslassungen von Suffixen (Suf-, Fehlerwert 1,74 pro 1.000 Tokens) oder Präfixen (Prä-, Fehlerwert 1,61 pro 1.000 Tokens) sowie überflüssige Präfixe (Prä+, Fehlerwert 1,36 pro 1.000 Tokens). Überflüssige Suffixe kamen im Korpus selten vor.

**Abbildung 29:** Fehlerverteilung bei der Affixschreibung; Fehlerwert pro 1.000 Tokens

Die Analyse der PoS-Tags der Zielhypothese ([ZHpos]) deutet darauf hin, dass die meisten verwechselten Suffixe (*Suf) bei den Substantiven (n = 128) geschahen. Besonders häufig sind hier substantivische Derivationen mit dem Suffix {-er} (z. B. *<Fernsehen> statt *Fernseher*; *<Bundesbürge> statt *Bundesbürger*). Dies demonstriert einerseits spezifische Schwierigkeiten der chinesischen Lernenden bei der Suffixderivation mit {-er}. Andererseits könnten solche Verwechslungen zwischen <e> und <er> auch auf lautlicher Ähnlichkeit beruhen, sodass die Lernenden beim Diktieren die Endungen nicht deutlich unterscheiden konnten, insbesondere wenn beide Formen als Substantive im Deutschen existieren (*Fernsehen* vs. *Fernseher*; *Lehren* vs. *Lehrer*). Bei Verben treten Abweichungen vor allem beim verbalen Suffix {-ier-} auf (z. B. *<motiwert>, *<motiwirt>, *<motivert> für *motiviert*; *<Studiren> für *Studieren*; *<informiren> für *informieren*). Die Fehlschreibung *<motiwert> und *<motiwirt> deutet darauf hin, dass etwa die Hälfte aller Graphemverwechslungen <w> für <v> (vgl. Kapitel 4.2.5) auf die Interaktion von morphologischer Stammschreibung und Suffixschreibung zurückzuführen sind. Zudem beweisen Verwechslungen {-ir-} für {-ier-}, dass die chinesischen Lernenden generelle Schwierigkeiten bei der Markierung der Vokalquantität aufweisen, unabhängig davon, ob diese im Stamm, in der Fuge oder im Affix auftritt.

Bei Adjektiven sind Abweichungen vor allem bei den Suffixen {-isch} (n = 24), {-ig} und {-lich} festzustellen. So wird {-isch} ganz häufig durch {-ich}, {-ig} oder {-lich} ersetzt, und auch {-ig} wird oft mit {-ich} oder {-isch} verwechselt. Im Deutschen gehören diese Suffixe zu den kategorieverändernden Suffixen, mit denen Substantive zu Adjektiven abgeleitet werden können (vgl. Eisenberg, 2013, S. 261), wie z. B. *Technik* zu *technisch*; *Stress* zu *stressig*, *Freund* zu *freundlich*. Die semantische und funktionale Differenzierung dieser Suffixe ist komplex. Adjektive mit {-ig} gehören zur Gruppe der Ornativa und beschreiben konkrete, physische oder stoffliche Eigenschaften einer Sache. Im Gegensatz dazu werden Adjektive mit {-isch} häufig von Basen abgeleitet, die Lebewesen bezeichnen, und drücken aus, dass etwas „in der Art von“ der Basis ist. Adjektive mit {-lich} hingegen sind funktional neutraler und beschreiben abstrakte, relationale oder zeitliche Eigenschaften (vgl. Eisenberg, 2013, S. 404; Fleischer & Barz, 2012, S. 343). Erschwerend kommt hinzu, dass von derselben Basis unterschiedliche Adjektive abgeleitet werden können (z. B. *geistlich* vs. *geistig*). Trotz der identischen Substantivbasis „der Geist“ haben *geistlich* und *geistig* unterschiedliche Bedeutungen. Während *geistlich* die Religion, den kirchlichen und gottesdienstlichen Bereich umfasst, bezieht sich *geistig* auf den menschlichen Geist und das Denkvermögen des Menschen. Diese Vielfalt und die subtilen Bedeutungsunterschiede machen die korrekte Suffixwahl zu einer großen Herausforderung. In der Studie von M. Liu (2019, S. 77) wurde beispielsweise die Abweichung aufgrund der Verwechslung zwischen den Adjektiven *geistig* und *geistlich* im Lernertext aufgedeckt: „Für mich sind beide körperliche und *geistliche Gesundheit ganz wichtig [...]“.

Die Verwechslung der Präfixe (*Prä) ist ebenfalls eine wichtige Fehlerquelle. Eine ebenenübergreifende Analyse mit [ZHpos] ergibt, dass diese Abweichungen hauptsächlich bei Verben (n = 71) vorkommen. Am häufigsten treten Fehler beim Präfix {-ver-} auf; oft wird stattdessen die Präposition *für* geschrieben (z. B. *<füreinsamen> für *vereinsa-*

men). Diese Beobachtung korreliert damit, dass über 82,26 % aller Fehlschreibungen <f> für <v> (n = 62) im phonographischen Bereich (vgl. Kapitel 4.2.5) auf Abweichungen des Präfixes {ver} (n = 51) zurückzuführen sind. Der hohe Fehlerwert bei {ver-} kann einerseits mit der zahlenmäßigen Dominanz von *ver*-Verben im Deutschen erklärt werden (vgl. Fleischer & Barz, 2012, S. 389). Andererseits ist die Präposition *für* den chinesischen Lernenden möglicherweise vertrauter als spezifische Präfixverben wie *vereinsamen*. In der Untersuchung von J. Wang (2017, S. 82) wird die Lernhürde im Umgang mit Affixen für chinesische Deutschlernende festgestellt, da sich die Wortbildungsverfahren chinesischer Verben und deutscher Präfixverben grundlegend unterscheiden. Chinesische Verben sind meistens zusammengesetzte Verbkomponenten, haben hingegen keine Präfixe.

Auch die Verwechslung von {b} mit dem Präfix {be-} (z. B. *<bleibt> statt *beliebt*) ist im DeChiLKo-Korpus festzustellen. Dies deutet einerseits auf Schwierigkeiten bei der Verschriftlichung des Schwa-Lauts im unbetonten Präfix hin. Andererseits korrespondieren solche Fehlschreibungen mit Ergebnissen vieler Studien zur Wortbetonung, die zeigen, dass chinesische Deutschlernende Schwierigkeiten haben, die Betonungsmuster im Deutschen korrekt wahrzunehmen und zu produzieren (vgl. Hunold, 2009, S. 156–166; Lay, 2017, S. 33; T. Liu, 2015, S. 84). Da im Deutschen nicht trennbare Verbpräfixe wie *be-*, *ent-*, *er-*, *ge-*, *ver-*, *zer-* keine Hauptbetonung tragen (vgl. Eisenberg, 2016, S. 42), können die Fehlschreibungen in Diktattexten eine mangelnde perzeptive Salienz dieser unbetonten Präfixe widerspiegeln.

Zusammenfassend lässt sich für den Bereich der Wortbildung festhalten, dass chinesische Deutschlernende sowohl bei der Kompositions- als auch bei der Affixschreibung deutliche Schwierigkeiten zeigten, wobei Fehler bei Affixen häufiger auftraten. Die Ursachen sind vielfältig und reichen von L1-Transfereffekten (Fehlen von Fugenelementen und Präfixen im Chinesischen) über die Komplexität und Ungeheimheiten der deutschen Wortbildung (Unregelmäßigkeit der Fugen, semantische Differenzierung von Suffixen) bis hin zu Problemen bei der Perzeption unbetonter Affixe.

4.5 Der syntaktische Bereich

Die syntaktische Ebene der deutschen Orthographie regelt Schreibkonventionen, die über die reine Wortebene hinausgehen und die Struktur von Sätzen sowie die Funktion von Wörtern im Satzkontext betreffen. Im Rahmen des DeChiLKo-Korpus fokussiert die Analyse in diesem Bereich auf einige der zentralen und für Lernende oft fehleranfälligen Phänomene der deutschen Rechtschreibung. Dazu zählen:

- die Groß- und Kleinschreibung ([Syn_GKS]): Abweichungen von den komplexen Regeln zur Majuskel- und Minuskelverwendung
- die Getrennt- und Zusammenschreibung ([Syn_GZS]): Fehler bei der korrekten Zusammensetzung bzw. Trennung von Wortbestandteilen

- die Unterscheidung von *man* und *Mann* ([Syn_man-Mann]): Verwechslungen zwischen dem Indefinitpronomen *man* und dem Substantiv *Mann*²⁵
- die Unterscheidung von *dass* und *das* [Syn_dass-das]: Abweichungen bei der Schreibung der Konjunktion *dass* und des Artikels/Pronomens *das*

Die Interpunktion, obwohl sie ebenfalls einen wichtigen Aspekt der syntaktischen Gliederung darstellt, wird in der vorliegenden Untersuchung nicht detailliert behandelt.²⁶

In den folgenden Unterkapiteln werden die spezifischen Fehlerphänomene in diesen vier Bereichen der syntaktischen Ebene detailliert dargestellt.

4.5.1 Fehler bei der Groß- und Kleinschreibung [Syn_GKS]

Als ein zentraler Teilbereich der deutschen Orthographie spielt die Groß- und Kleinschreibung (GKS) im Spracherwerb eine wichtige Rolle. Auch für die Lernenden mit Deutsch als Erstsprache stellt der korrekte Umgang mit der deutschen Großschreibung eine anhaltende Herausforderung dar (vgl. H. Günther & Gaebert, 2011; Röber, 2011, S. 296) und Fehler in diesem Bereich gelten nach wie vor als eine signifikante orthographische Fehlerquelle im Abitur (vgl. Fuhrhop & Romstadt, 2021). Im DaF-/DaZ-Kontext bestätigen viele Studien, dass die Groß- und Kleinschreibung ein erhebliches Lernhindernis für die Lernenden mit unterschiedlichen L1 darstellt (vgl. Fix, 2002; Maas, 2016; Nimz, 2022; Steinig et al., 2009; Tang, 2003; Thomé, 1987; Z. Wu & Li, 2022).

Auf der Annotationsebene [Syn_GKS] werden die Abweichungen primär zwischen Groß- für Kleinschreibung (*GfK*) und Klein- für Großschreibung (*KfG*) unterschieden. Innerhalb der Kategorie Klein- für Großschreibung erfolgt eine weitere Differenzierung in Unterkategorien: Kleinschreibung von Substantiven (*KfG_Sub*), von Satzanfängen (*KfG_SA*), der Eigennamen (*KfG_Eigennamen*), von Anredepronomen (*KfG_Anrede*) sowie sonstige Fälle von Klein- für Großschreibung (*KfG_Son*). Unter den Abweichungen Groß- für Kleinschreibung (*GfK*) wird zusätzlich die Binnengroßschreibung (Großschreibung im Wort) mit dem Tag *GiW* markiert.

Insgesamt wurden 604 Abweichungen im Bereich der Groß- und Kleinschreibung in 209 Lernertexten im DeChiLKo-Korpus gefunden. Der Fehlerwert pro 1.000 Tokens beträgt 19,07 (vgl. Tabelle 36).

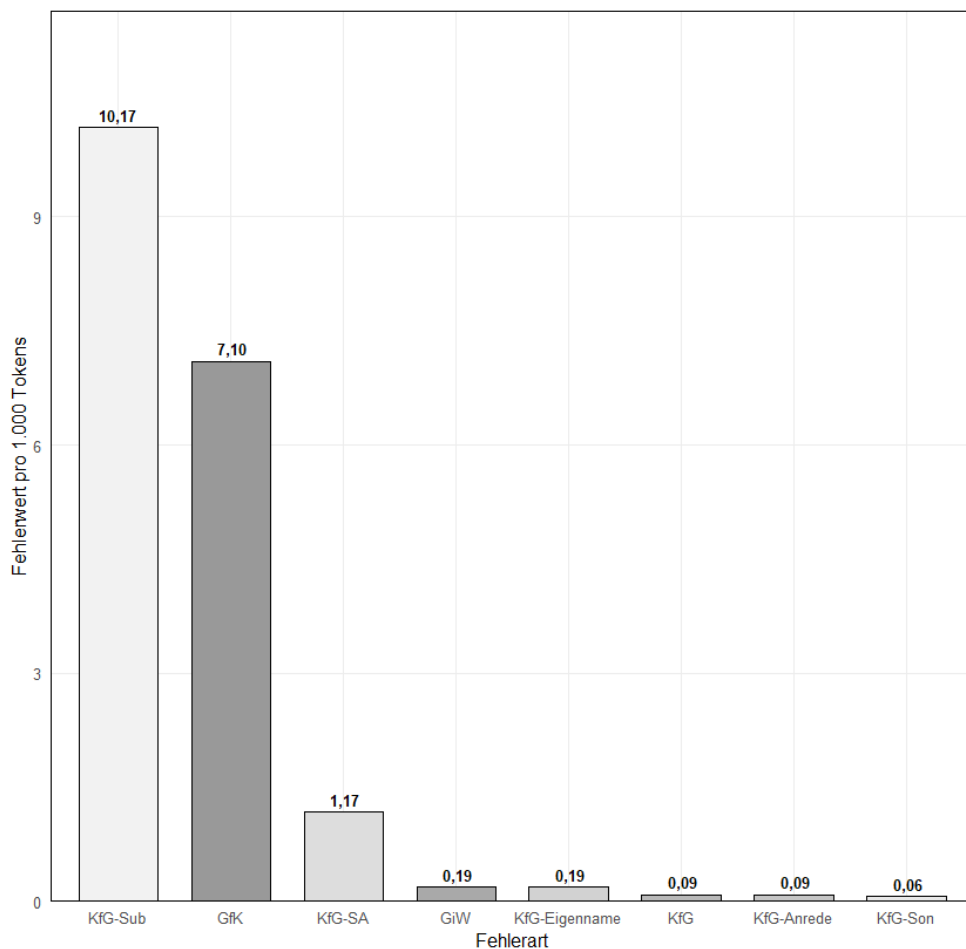
25 Obwohl es sich bei *man* und *Mann* streng genommen um unterschiedliche Lexeme handelt und eine Verwechslung somit auch als lexikalisch-morphologischer Fehler interpretiert werden könnte, wird dieses Phänomen im Annotationsschema der vorliegenden Arbeit bewusst der syntaktischen Ebene zugeordnet. Der Grund hierfür ist, dass die korrekte Verschriftlichung zwingend die Identifikation der syntaktischen Funktion (Pronominalgebrauch vs. nominaler Nukleus) im jeweiligen Satzkontext voraussetzt.

26 Die Erhebung der Lernerdaten im DeChiLKo-Korpus erfolgte mittels Diktaten. In dieser Aufgabenform werden Satzzeichen explizit angesagt, um den Diktatfluss zu steuern und eine normgerechte Textwiedergabe zu ermöglichen. Eine Analyse der Interpunktionssetzung in den Diktattexten würde daher nicht die tatsächliche Interpunktionskompetenz der Lernenden widerspiegeln, sondern primär deren Fähigkeit, angesagte Satzzeichen korrekt zu notieren. Aus diesem Grund wird auf eine gesonderte Analyse der Interpunktion im Rahmen dieser Arbeit verzichtet.

Tabelle 36: Überblick über die Abweichungen der Groß- und Kleinschreibung

Abweichungen der Groß- und Kleinschreibung insgesamt	Summe aller Tokens	Abweichungen der Groß- und Kleinschreibung pro 1.000 Tokens
604	31.674	19,07

Die folgende Abbildung 30 zeigt die Fehlerverteilung je nach Fehlerart in dem Bereich Groß- und Kleinschreibung.

**Abbildung 30:** Fehlerverteilung bei der Groß- und Kleinschreibung; Fehlerwert pro 1.000 Tokens

Aus der Fehlerverteilung ergibt sich, dass Abweichungen bei der Großschreibung häufiger als Fehler bei der Kleinschreibung erscheinen. Die mit Abstand wichtigste Fehlerquelle in den Lernertexten ist die Kleinschreibung von Substantiven (KfG_Sub), deren Fehlerwert 10,17 pro 1.000 Tokens beträgt. Darauf folgt die Großschreibung statt Klein-

schreibung (GfK) mit einem Fehlerwert von 7,10 pro 1.000 Tokens. Mit deutlichem Abstand rangiert die Kleinschreibung am Satzanfang (KfG_SA) mit einem Fehlerwert von 1,16 pro 1.000 Tokens an dritter Stelle. Andere Fehlerarten der Groß- und Kleinschreibung, wie z. B. bei Eigennamen, Anredepronomen oder die Binnengroßschreibung, treten im Korpus nur selten auf.

Eine genauere Betrachtung der am häufigsten kleingeschriebenen Substantive im Korpus erfolgt in der folgenden Tabelle 37:

Tabelle 37: Die am häufigsten kleingeschriebenen Substantive im Korpus

ZH (Zielwort)	Anzahl	Fehlerwert pro 1.000 Tokens
Weiterlernen	26	0,82
Lerndauer	19	0,60
Stress	14	0,44
Möglichkeiten	12	0,38
Anfang	12	0,38
Lehrer	11	0,35
Kinder	10	0,32
Belohnungen	10	0,32
Zeit	10	0,32

Deutlich kann man an der Tabelle erkennen, dass das substantivierte Verb *Weiterlernen* am häufigsten betroffen ist. Beim Wort *Weiterlernen* handelt es sich um eine Konversion, bei der die Wortart vom Infinitiv *weiterlernen* zum Substantiv *Weiterlernen* wechselt, ohne dass sich die Wortform ändert. Gemäß dem Amtlichen Regelwerk (§ 57, Geschäftsstelle des Rats für deutsche Rechtschreibung [Hrsg.], 2024) werden Wörter anderer Wortarten großgeschrieben, wenn sie als Substantive gebraucht werden (Substantivierungen). Ein wichtiges Merkmal für den Substantivgebrauch ist der vorausgehende Artikel oder eine Präposition mit Artikel (vgl. Fuhrhop, 2021, S. 41). Obwohl das substantivierte Verb *Weiterlernen* im Diktattext in einer Präpositionalphrase *zum Weiterlernen* deutlich als nominale Einheit signalisiert und von den meisten Lernenden auch mit der Verschmelzung *zum* geschrieben wurde, wurde die Großschreibung dennoch häufig vernachlässigt.

Unter den häufig kleingeschriebenen Substantiven sind viele Abstrakta zu finden (z. B. *Lerndauer*, *Stress*, *Anfang*, *Zeit*). Diese Ergebnisse korrespondieren mit den Befunden von Liu (2019) zu den Fehlererscheinungen in Aufsätzen chinesischer Deutschlerner am Gymnasium Ettal, wo sich ebenfalls die meisten kleingeschriebenen Nomen der Gruppe der Abstrakta zuordnen ließen (vgl. Liu, 2019, S. 48–49). Darüber hinaus treten Kleinschreibungen auch häufig bei abgeleiteten Nomen auf, die mithilfe typischer nominalisierender Suffixe gebildet werden (z. B. *Möglichkeiten*, *Lehrer*, *Belohnun-*

gen). Solche Fehlschreibungen deuten darauf hin, dass die chinesischen Lernenden die nominalisierende Funktion wichtiger Suffixe wie {-keit}, {-ung}, {-er} noch nicht verinnerlicht haben oder diese nicht konsistent als Signal für eine Substantivierung und somit Großschreibung erkennen.

In der Kategorie Großschreibung für Kleinschreibung (GfK) sind fehlerhafte Großschreibungen der Adjektive (n = 56) häufig. Bemerkenswert ist hierbei, dass die meisten der fälschlicherweise großgeschriebenen Adjektive in ihrer deklinierten Form mit Endungen stehen (z. B. *technischen*, *junge*, *deutsche*).

Bei der fehlerhaften Großschreibung von Verben (n = 81) sind vor allem Infinitive (n = 33) betroffen, zu denen es homophone Substantive gibt (z. B. *fahren* (das *Fahren*), *wünschen* (der *Wunsch*), *Kochen* (das *Kochen*), *kosten* (die *Kosten*)).

Die festgestellte Dominanz der Kleinschreibung von Substantiven als Hauptfehlerquelle bei chinesischen Deutschlernenden ist im Kontext ihres multi-kompetenten Repertoires zu sehen. Ein Vergleich mit den Ergebnissen in Thomés Studie (vgl. 1987, S. 102) macht deutlich, dass Lernende mit Türkisch als L1 im Bereich der Groß- und Kleinschreibung einen Fehlerwert von 15,34 pro 500 Wörter (entspricht 30,68/1.000 Wörter) aufweisen. Im Gegensatz dazu erreichten die chinesischen Deutschlernenden im DeChiLko-Korpus mit einem Gesamtwert von 19,07 Fehlern pro 1.000 Tokens für alle GKS-Fehler eine tendenziell geringere Fehleranfälligkeit in diesem spezifischen Bereich. Obwohl die L1 Chinesisch keine alphabetische Sprache ist und somit auch keine Unterscheidung zwischen Groß- und Kleinbuchstaben kennt, legt die chinesische Schrift als eine primär logographische Zeichensprache großen Wert auf die exakte visuelle Form der Schriftzeichen. Es ist denkbar, dass diese in der L1 geschulte Aufmerksamkeit für die äußere Form sprachlicher Einheiten von den Lernenden als Lernstrategie auf den Erwerb der deutschen Orthographie übertragen wird, und daher konzentrieren sie sich eher auf die formalen Aspekte der Groß- und Kleinschreibung. Die dennoch auftretenden Schwierigkeiten, insbesondere bei der Großschreibung von Substantivierungen (z. B. *Weiterlernen*) und Abstrakta, deuten jedoch darauf hin, dass die chinesischen Lernenden die morpho-syntaktischen Signale, die eine Substantivierung im Deutschen ausmachen (z. B. Artikel, Präpositionen, nominalisierende Suffixe), noch nicht sicher erkennen.

Die fehlerhafte Großschreibung von deklinierten Adjektiven könnte auf eine Übergeneralisierung zurückzuführen sein. Für die meisten chinesischen Lernenden ist die Kleinschreibung der Adjektive im Deutschen bekannt, weil ähnliche Regeln in ihrer ersten Fremdsprache Englisch existieren. Jedoch kann ein Adjektiv im Deutschen durch Konversion auch ein Substantiv mit gleicher Stammform sein, wie z. B. *jung* – *der Junge*, *gut* – *der/die/das Gute*. Phonetisch sind diese Formen in der gesprochenen Sprache identisch (vgl. Eisenberg, 2013, 279 f.). Die Lernenden scheinen beim Diktat Schwierigkeiten zu haben, die syntaktische Funktion des Wortes im Satzgefüge rechtzeitig zu analysieren, weshalb sie die ambigen Formen fälschlicherweise oft als Substantivierungen interpretieren und großschreiben.

Die fehlerhafte Großschreibung von Länderadjektiven (z. B. *<Deutsche> statt <deutsche>) ist ein Hinweis auf einen Transfer aus der L2 Englisch. Im Englischen

werden Adjektive, die sich auf Nationalitäten beziehen, großgeschrieben (*German literature*). Im Gegensatz dazu schreibt man in der deutschen Sprache adjektivische Ableitungen von Eigennamen wie *deutsch*, *chinesisch* klein (vgl. Hrsg., 2024, S. 98), es sei denn, sie fungieren als fester Bestandteil eines mehrteiligen Eigennamens (wie z. B. die *Hamburger Schreibprobe* oder der *Zweite Weltkrieg*). Es ist daher anzunehmen, dass die chinesischen Lernenden die ihnen aus dem Englischen vertraute Großschreibungsregel für Nationalitätenadjektive auf das Deutsche übertragen.

Die fälschliche Großschreibung von Verben, insbesondere bei Infinitiven mit homophonen Substantiven, deutet ebenfalls auf Schwierigkeiten bei der Unterscheidung von Wortarten im Kontext hin. Da der einfache Infinitiv und das konvertierte Substantiv phonologisch völlig deckungsgleich sind, erfordert die korrekte orthographische Realisierung eine tiefergehende syntaktische Analyse während des Schreibprozesses (z. B. das Erkennen von vorausgehenden Artikeln oder Präpositionen als Substantivierungssignale). Diese kognitive Transferleistung scheint für die Lernenden in der Echtzeit-Situation des Diktats eine erhebliche Hürde darzustellen.

Deutlich seltener sind Fehler bei der Kleinschreibung am Satzanfang, weshalb sie nur eine untergeordnete Fehlerquelle repräsentieren. Der niedrige Fehlerwert in diesem Bereich weist darauf hin, dass die satzinitiale Großschreibung für die chinesischen Lernenden im Allgemeinen kein Problem darstellt. Dies ist einerseits darauf zurückzuführen, dass die Erkennung eines Satzanfangs im Vergleich zur Bestimmung eines Substantivs oder der korrekten Schreibung von Adjektiven oft einfacher ist. Andererseits sind die chinesischen Lernenden durch den Erwerb ihrer L2 Englisch bereits mit der satzinitialen Großschreibung vertraut, was einen positiven Transfer auf das Deutsche begünstigt. Darüber hinaus existiert die Regel der Großschreibung am Satzanfang auch für *Pinyin*-Sätze (vgl. Offizielle Norm vom Ministry of Education der VR China, 29.2012). Obwohl vollständige in *Pinyin* geschriebene Sätze selten in der alltäglichen Schriftpraxis vorkommen, begegnet man solchen Sätzen, die mit einem Großbuchstaben beginnen, beispielsweise in Grundschulbüchern oder als Aussprachehilfe. Dieses parallele Regelbewusstsein könnte zur besseren Beherrschung der satzinitialen Großschreibung im Deutschen beitragen. Es ist jedoch zu berücksichtigen, dass alle Satzzeichen in der Diktataufgabe explizit vorgelesen werden, was den Lernenden die Identifizierung des Satzanfangs zusätzlich erleichtern dürfte.

Die Analyse der Groß- und Kleinschreibungsfehler im DeChiLKo-Korpus offenbart eine komplexe Vernetzung verschiedener Faktoren innerhalb des multi-kompetenten Systems der chinesischen Deutschlernenden. Die Kleinschreibung von Substantiven stellt hierbei die grundlegende Herausforderung und den dominantesten Fehlertyp dar. Damit verbundene Unsicherheiten bei Substantivierungen und Abstrakta deuten auf Schwierigkeiten hin, die relevanten morpho-syntaktischen Signale der deutschen Sprache korrekt zu interpretieren. Darüber hinaus sind auch Transfereffekte aus der zuvor erworbenen englischen Sprache in diesem Bereich zu beobachten, die sich sowohl positiv (z. B. bei der Satzanfangsgroßschreibung) als auch negativ (z. B. bei Nationalitäten-

adjektiven) auswirken können. Schließlich tragen die inhärente Komplexität der deutschen Nominalisierungsprozesse sowie die oft kontextabhängige und im Rahmen einer Diktataufgabe zusätzlich erschwerte Wortarterkennung zu den Fehlschreibungen bei.

4.5.2 Fehler bei der Getrennt- und Zusammenschreibung [Syn_GZS]

Neben der Groß- und Kleinschreibung stellt auch die Getrennt- und Zusammenschreibung (GZS) für viele Lernende des Deutschen eine Herausforderung dar, sowohl für L1-Sprechende als auch im DaF/DaZ-Kontext. Studien wie die von Nimz (vgl. 2022, S. 44) weisen auf signifikante Unterschiede bei syntaktischen Schreibungen (GKS und GZS) zwischen bilingualen und monolingualen Schüler*innen hin.

Auf der Annotationsebene [Syn_GZS] werden in DeChiLKo folgende Fehlerkategorien unterschieden:

- Getrennt- statt Zusammenschreibung (*GfZ*): Wörter, die fälschlicherweise getrennt geschrieben werden.
- Zusammen- statt Getrenntschreibung (*ZfG*): Wörter, die fälschlicherweise zusammengeschrieben werden.
- Getrenntschreibung von unselbstständigen Teilen (*GUT*): Fehlerhafte Trennung von Wortteilen, die keine eigenständigen Einheiten bilden (z. B. Affixe vom Stamm).
- Abweichungen bei der Anwendung von Divis (*Divis*): Fehlerhafte oder fehlende Verwendung des Bindestrichs.

Insgesamt wurden 731 Abweichungen im Bereich der Getrennt- und Zusammenschreibung in 255 Lernertexten im DeChiLKo-Korpus gefunden. Der Fehlerwert pro 1.000 Tokens beträgt 23,08 (vgl. Tabelle 38).

Tabelle 38: Überblick über Abweichungen der Getrennt- und Zusammenschreibung

Abweichungen der Getrennt- und Zusammenschreibung insgesamt	Summe aller Tokens	Abweichungen der Getrennt- und Zusammenschreibung pro 1.000 Tokens
731	31.674	23,08

Die folgende Abbildung stellt die genaue Fehlerverteilung unterschiedlicher Fehlerarten im Bereich der Getrennt- und Zusammenschreibung dar.

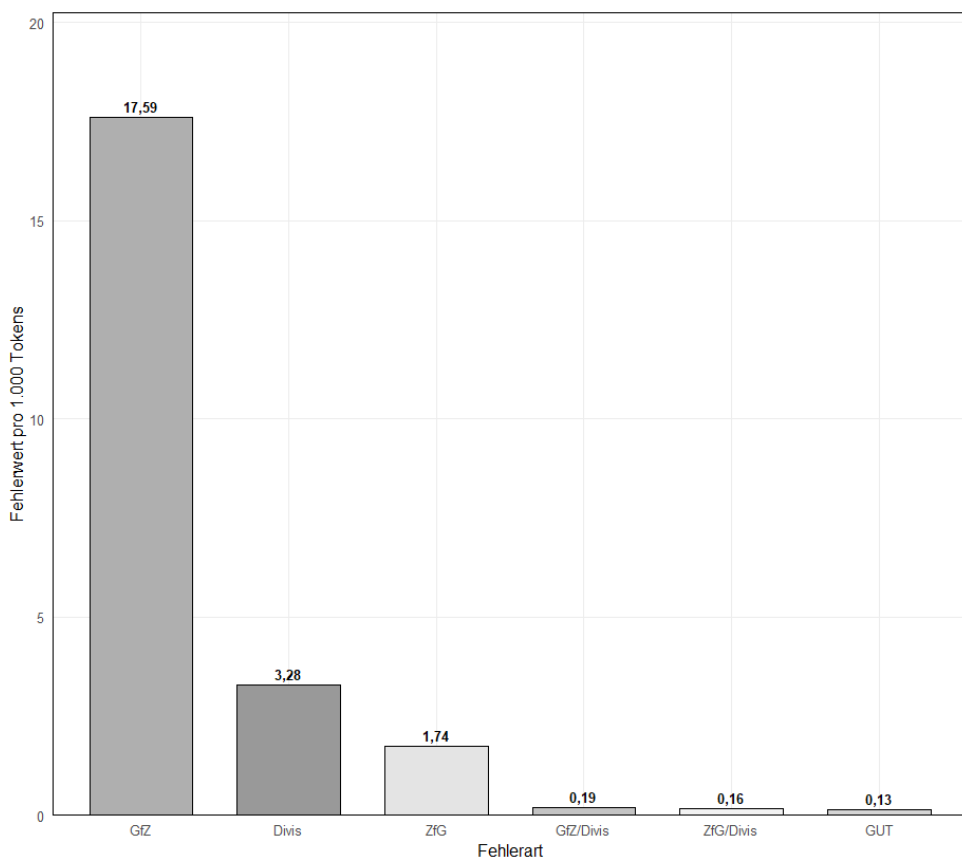


Abbildung 31: Fehlerverteilung der Getrennt- und Zusammenschreibung; Fehlerwert pro 1.000 Tokens

Die Fehlerverteilung (vgl. Abbildung 31) zeigt deutlich, dass die Getrenntschreibung statt Zusammenschreibung (*GfZ*) mit einem Fehlerwert von 17,58 pro 1.000 Tokens die wichtigste Fehlerquelle in dieser Kategorie ist. Mit einem großen Abstand folgt die fehlerhafte Anwendung von *Divis* (*Divis*) mit einem Fehlerwert von 3,27 pro 1.000 Tokens. An dritter Stelle steht die Zusammenschreibung statt Getrenntschreibung (*ZfG*). Die Getrenntschreibung unselbstständiger Teile (*GUT*) sowie das Zusammentreffen mehrerer Fehlerarten (*GfZ/Divis*; *ZfG/Divis*) kommen im Korpus vergleichsweise selten vor.

Unter der Kategorie „Getrenntschreibung statt Zusammenschreibung“ (*GfZ*) werden Substantive ($n = 298$) am häufigsten fehlerhaft getrennt geschrieben. Betroffen sind hier vor allem Zusammensetzungen wie Freizeitangebote ($n = 76$), Freizeitvergnügen ($n = 80$), Weiterlernen ($n = 69$), Lehrinhalte ($n = 79$) sowie Lerndauer ($n = 58$). Eine Gemeinsamkeit dieser Beispiele ist das Fehlen eines Fugenelements zwischen den Gliedern, was die Erkennung als Kompositum für die Lernenden möglicherweise erschwert. Zudem handelt es sich bei *Freizeitvergnügen* und *Weiterlernen* um Zusam-

mensetzungen mit substantivierten Verben als Endglied. Es ist denkbar, dass die Lernenden beim Diktat die Substantivierung fehlerhaft als Verb interpretierten und sie daher getrennt von den anderen Wortteilen schrieben.

Auch nicht zusammengesetzte mehrteilige Adverbien ($n = 90$) werden häufig getrennt geschrieben (z. B. *<einer seits> für <einerseits>; *<anderer seits> für <andererseits>; *<vor her> für <vorher>). Der Grund hierfür könnte darin liegen, dass die Erstglieder dieser Adverbien oft auch als eigenständige lexikalische Elemente vorkommen und die Lernenden sie daher fälschlicherweise getrennt schrieben.

Bei Verben sind vor allem diejenigen mit untrennbaren Präfixen auffällig, die getrennt geschrieben werden (z. B. *vereinsamen*, *veranschaulichen*). Diese Getrenntschreibungen stehen jedoch meist im Zusammenhang mit den bereits diskutierten Verwechslungen des Präfixes {ver-} mit der Präposition *für* (vgl. Kapitel 4.4.2 zur Affixschreibung). Fälle, in denen die Lernenden unselbstständige Präfixe tatsächlich isoliert vom Verbstamm getrennt schrieben, sind im Korpus sehr selten. Dies deutet darauf hin, dass die chinesischen Lernenden bereits eine grundlegende Kompetenz im Bereich der Zusammenschreibung von (nicht trennbaren) Präfixen und Verbstämmen entwickelt haben. Auffälliger ist hingegen die fehlerhafte Getrenntschreibung der Verbpartikel bei trennbaren Verben in Infinitiv-, Partizipial- oder Nebensatzkonstruktionen, bei denen das Verb am Ende steht (z. B. *<aus schlafen> statt <ausschlafen>; *<fest legt> statt <festlegt>; *<aus ruhen> statt <ausruhen>; *<zurück gekommen> statt <zurückgekommen>). Besonders häufig tritt diese Getrenntschreibung bei Verben auf, deren Partikeln auch als eigenständige Präpositionen oder Adverbien existieren. Erschwerend kommen Zweifelsfälle in der deutschen Rechtschreibung hinzu, wie z. B. beim Verb *kennenlernen*, bei dem heutzutage sowohl die Zusammen- als auch die Getrenntschreibung zulässig sind²⁷. Gleiches gilt für Verben wie *sitzen bleiben* (*sitzenbleiben*) oder *stehen lassen* (*stehenlassen*) in übertragener Bedeutung (vgl. Geschäftsstelle des Rats für deutsche Rechtschreibung [Hrsg.], 2024, S. 59). Solche Unklarheiten und Zweifelsfälle bei der Getrennt- und Zusammenschreibung von Verben stellen für Lernende eine zusätzliche Hürde dar.

Alle Abweichungen hinsichtlich der Divis-Schreibung betreffen im DeChiLKo-Korpus das Auslassen des Ergänzungsstrichs. Am häufigsten kommt die Auslassung des Divises bei Wortgruppen wie *Spiel- und Lerndauer* ($n = 86$) vor, wo das als Ergänzungsstrich fungiert. Der Ergänzungsstrich ersetzt hier einen wiederholten Wortteil (*Spieldauer und Lerndauer* → *Spiel- und Lerndauer*) und dient der besseren Lesbarkeit sowie der Vermeidung von Wiederholungen. Der hohe Fehlerwert bekundet, dass die Lernenden diese Funktion des Divises oft nicht erkennen. Dies könnte darauf zurückzuführen sein, dass die Anwendung des Ergänzungsstrichs weder in ihrer L1 Chinesisch noch in ihrer ersten Fremdsprache Englisch eine direkte Entsprechung hat. Das Englische und das Chinesische neigen dazu, Wiederholungen durch parallele Konstruktionen zu vermeiden, ohne einen Strich einzusetzen, z. B. für die Wortgruppe *Geschäfts- und Privatkunden* wird im Englischen als *business and private customers* und im

27 Im DeChiLKo wurden getrennt geschriebene *kennenlernen* jedoch als Verstöße gegen die Regel der Zusammenschreibung annotiert, weil es in der Zielhypothese (Diktattext) als zusammengesetztes Verb geschrieben wird.

Chinesischen als 商业与私人客户 (shāngyè yǔ sīrén kèhù, wörtl.: Geschäft und privat Kunden) formuliert. Des Weiteren wird der Divis als Bindestrich auch häufig bei Koppelungen ausgelassen, wie z. B. in **Dreizimmer Wohnung* (3-Zimmer-Wohnung), **Tischtennis Verein* (Tischtennis-Verein) oder **Humboldt Universität* (Humboldt-Universität). Bei mehrteiligen Zusammensetzungen wie 3-Zimmer-Wohnung ist der Bindestrich obligatorisch, um die Verbindung zwischen der Ziffer und dem Nomen zu verdeutlichen, wobei Bindestriche zwischen allen Bestandteilen gesetzt werden (§ 44, Geschäftsstelle des Rats für deutsche Rechtschreibung [Hrsg.], 2024, S. 71). Bei Zusammensetzungen wie *Tischtennis-Verein* dient der Bindestrich primär der besseren Lesbarkeit und Strukturierung. Enthält eine Zusammensetzung Eigennamen als erste Bestandteile, wie in *Humboldt-Universität*, wird nach § 50 ein Bindestrich gesetzt (vgl. Geschäftsstelle des Rats für deutsche Rechtschreibung [Hrsg.], 2024, S. 75). Schreibweisen wie <Fitnessstudio> oder <Fitness Studio> (*Fitness-Studio*) im Korpus wurden zunächst als Auslassungen des Bindestrichs markiert, sind jedoch nach §E1 (vgl. Geschäftsstelle des Rats für deutsche Rechtschreibung [Hrsg.], 2024, S. 73) orthographisch nicht als Fehler zu werten. Solche aus dem Englischen stammenden Verbindungen aus Substantiv + Substantiv wie (*Fitness + Studio*) werden aufgrund ihrer grammatischen Funktion als Zusammensetzung in der Regel zusammengeschrieben (*Fitnessstudio*); als Variante ist jedoch auch die verdeutlichende Schreibung mit Bindestrich (*Fitness-Studio*) möglich. Die Komplexität und Uneinheitlichkeit der Regeln zur Bindestrichverwendung im Deutschen könnten das Lernen und die Anwendung dieser orthographischen Konventionen erschweren. Besonders für chinesische Lernende, in deren Erstsprache keine vergleichbare Funktion des Bindestrichs existiert, könnte dies eine Herausforderung darstellen.

Unter der Kategorie „Zusammenschreibung für Getrenntschreibung“ (ZfG) treten häufig Fehler bei nominalen Phrasen auf, die fälschlicherweise zusammengeschrieben werden, wie z. B. **Technischemöglichkeit* (*technische Möglichkeit*) **Jungeleute* (*junge Leute*), **Großenappetit* (*großen Appetit*), **Deutschsprache* (in der *deutschen Sprache*). Solche Fehlschreibungen lassen sich auf eine Übergeneralisierung der deutschen Kompositabilität zurückführen. Die Lernenden interpretierten nominale Phrasen fälschlicherweise als Komposita und wandten entsprechend das Prinzip der Zusammenschreibung von Wortverbindungen mit zwei oder mehr Stämmen an. Die Fehlschreibung **Technischemöglichkeit* deutet beispielsweise darauf hin, dass die Lernenden den morphologischen Aufbau von Wortstämmen noch nicht vollständig verstanden haben und die Regel der Kompositabildung nicht sicher anwenden können. Fehler kommen auch bei der Schreibung von Infinitiven mit *zu* vor (z. B. **zufinden* statt *zu finden*). Darüber hinaus entstehen fehlerhafte Zusammenschreibungen durch die Verwechslung ähnlicher Ausdrücke, wie etwa *zufolge* statt *zur Folge*. Diese Fehlschreibungen deuten darauf hin, dass die Lernenden Schwierigkeiten haben, die funktionale und grammatische Rolle von *zu* in unterschiedlichen Kontexten korrekt zu analysieren.

Die Aufklärung der Fehler bei der Getrennt- und Zusammenschreibung signalisiert, dass chinesische Deutschlernende erhebliche Schwierigkeiten mit diesem komplexen Bereich der deutschen Orthographie aufwiesen. Die Dominanz der Getrenntschreibung statt Zusammenschreibung (GfZ) ist hierbei besonders auffällig. Im

Kontext der Multi-Kompetenz-Hypothese lassen sich die Fehlschreibungen als Ergebnis eines komplexen Zusammenspiels des gesamten sprachlichen Repertoires der chinesischen Lernenden interpretieren. Die L1 Chinesisch ist eine Sprache, in der Wörter in der Regel als isolierte Einheiten (Schriftzeichen) dargestellt werden, ohne dass es explizite Fugenelemente oder komplexe Zusammenschreibungsregeln gibt, wie sie im Deutschen üblich sind. Obwohl Komposita im Chinesischen häufig sind, werden sie meist durch Juxtaposition gebildet, nämlich durch das einfache Nebeneinanderstellen der Komponenten (vgl. M. Reichardt & Reichardt, 1990, S. 58). Diese L1-Prägung könnte zu einer Tendenz führen, Wortgrenzen auch im Deutschen stärker zu segmentieren. Die bereits erworbene Fremdsprache Englisch weist zwar ebenfalls Kompositionsmuster auf, die oft zusammengeschrieben werden (z. B. *homework*), es gibt auch viele Verbindungen, die getrennt bleiben (z. B. *high school*), oder solche, die mit Bindestrich verbunden werden (z. B. *ice-cream*). Aufgrund dieser unterschiedlichen und teilweise von den deutschen Regeln abweichenden Konventionen im Englischen ist eine direkte Übertragung von Wortbildungs- und Schreibprinzipien auf das Deutsche nicht möglich. Dies könnte sogar zu zusätzlichen Unsicherheiten beitragen.

Die Schwierigkeiten bei der Erkennung von Komposita ohne Fugenelement sowie von Substantivierungen als Kompositionsglieder deuten auf ein noch nicht gefestigtes Verständnis der deutschen Wortbildungsprinzipien hin.

Die fehlerhafte Getrenntschreibung von Verbpartikeln bei trennbaren Verben spiegelt die Komplexität der deutschen Verbstellung und Morphologie wider. Die Probleme mit dem Divis, insbesondere dem Ergänzungsstrich, sind ebenfalls durch den Transfereffekt geprägt, da diese spezifische Funktion des Bindestrichs weder im Chinesischen noch im Englischen bekannt ist. Die Komplexität und teilweise Optionalität der Bindestrichsetzung im Deutschen stellen zusätzliche Hindernisse dar. Die Zusammenschreibung statt Getrenntschreibung (ZfG), oft bei nominalen Phrasen, zeigt eine Übergeneralisierung des im Deutschen sehr produktiven Prinzips der Kompositabilisierung. Den Lernenden fällt es offenbar noch schwer, die Grenzen zwischen einer festen Zusammensetzung und einer losen syntaktischen Gruppe sicher zu ziehen.

4.5.3 Verwechslung zwischen *man* und *Mann* [Syn_*man-Mann*] sowie *dass* und *das* [Syn_*dass-das*]

Im Deutschen ist für das Leseverständnis die Unterscheidung der homophonen Wortpaare *das/dass* sowie *man/Mann* von entscheidender Bedeutung. Obwohl diese Paare denselben phonologischen Ausdruck (/das/ bzw. /man/) aufweisen, unterscheiden sie sich grundlegend in ihren syntaktischen Funktionen und ihrer Wortart (vgl. Fuhrhop & Peters, 2013, S. 242). Insbesondere die *das/dass*-Schreibung ist als eine persistente orthographische Fehlerquelle sowohl im Erstspracherwerb als auch im DaZ-Kontext bekannt. In der von Pießnack und Schübel (vgl. 2005, S. 67) durchgeführten Untersuchung von Abiturientenaufsätzen rangierte die *das/dass*-Schreibung an erster Stelle der fehlergefährdetsten Wörter. Auch in der Studie von Thomé (vgl. 1987, S. 137–138) waren die meisten Fehlschreibungen von den Graphemen <s> und <ß> bei türkischen Schü-

ler*innen auf der Verwechslung zwischen der Konjunktion (damals als *daß* geschrieben) und dem Artikel/Pronomen *das* zurückzuführen.

Auf der Annotationsebene [Syn_dass-das] werden im DeChiLKo-Korpus die Verwechslungen zwischen der Konjunktion *dass* sowie dem Artikel und dem Demonstrativpronomen *das* markiert. Auf der Annotationsebene [Syn_Mann-man] werden entsprechend die Verwechslungen zwischen dem Indefinitpronomen *man* sowie dem Substantiv *Mann* betrachtet.

Insgesamt werden 51 Abweichungen bei der Unterscheidung zwischen *dass* und *das* in 40 Lernertexten des gesamten Korpus festgestellt. Der Fehlerwert beträgt hier 1,61 pro 1.000 Tokens. Die Verwechslungen zwischen *Mann* und *man* traten hingegen lediglich 14-mal im Korpus auf. Da weder das Substantiv *Mann* noch das Pronomen *man* im Diktattext des Prüfungskorpus_2019 vorkamen, wurden bei der Berechnung des Fehlerwerts für dieses Paar ausschließlich das Erwerbskorpus sowie das Prüfungskorpus_2017 berücksichtigt. Der resultierende Fehlerwert pro 1.000 Tokens beträgt 0,63 (vgl. Tabelle 39).

Tabelle 39: Überblick über die Verwechslungen zwischen *dass* und *das* sowie *Mann* und *man*

Abweichungen	Anzahl	Summe aller Tokens	Abweichungen pro 1.000 Tokens
dass – das	51	31.674	1,61
Mann – man	14	22.295	0,63

Bei der *dass/das*-Schreibung sind die meisten Fehler auf die Falschschreibung der Konjunktion *dass* als *das* (n = 41) zurückzuführen. Die korrekte *dass*-Schreibung wird in der Orthographieerwerbsforschung oft nicht als isoliertes Wortproblem, sondern im Zusammenhang mit der Kommasetzung thematisiert (vgl. Fuhrhop, 2021, S. 87). Verschiedene Untersuchungen haben einen Zusammenhang zwischen der falschen Schreibung von *dass* und dem Fehlen von Kommata in Konjunktionalsätzen festgestellt (vgl. Feilke, 2011, S. 348–349). Obwohl im Rahmen der Diktataufgaben im DeChiLKo alle Satzzeichen (Punkt, Komma, Anführungszeichen usw.) explizit vorgelesen wurden, tritt im Korpus dennoch zehnmal das Zusammentreffen von Nichtkommatierung und der Falschschreibung **das* (statt *dass*) auf (siehe Beispiel 1).

1. [word]: Viele Eltern haben zwar die Sorge das kindern von dem Computer vereinsammeln.
[ZH]: Viele Eltern haben zwar die Sorge, dass Kinder vor dem Computer vereinsamen.
(DeChiLKo_Prüfung2017: C_NW_2017_D_002 (tokens 94–116))

Wesentlich häufiger (n = 31) kommen jedoch die Fehlschreibungen **das* statt *dass* trotz einer korrekten Kommasetzung vor (siehe Beispiel 2 und 3).

2. [word]: Ich glaube, das sie eine gute Täterin wird.
[ZH]: Ich glaube, dass sie eine gute Tänzerin wird.
(DeChiLKo_Erwerb: LZX20221022 (tokens 142–163))

3. [word]: Viele Eltern haben zwar die Sorge, dass Kinder vor dem Computer einsam.
 [ZH]: Viele Eltern haben zwar die Sorge, dass Kinder vor dem Computer vereinsamen.
 (DeChiLKO_Erwerb: C_N_2017_D_009 (tokens 84–106))

Diese Fehlschreibungen deuten darauf hin, dass die Lernenden die Verbindung zwischen der korrekten *dass*-Schreibung und den relevanten syntaktischen Merkmalen – wie der Verbletzstellung im Nebensatz sowie der obligatorischen Kommasetzung – nicht sicher erkennen oder anwenden. Nach Müller (2007, S. 70) wird „eine Kommasetzung umso wahrscheinlicher, je semantisch-funktional autonomer die vom Komma getrennten Äußerungseinheiten zueinander sind“. Der *dass*-Satz weist jedoch eine starke syntaktische Bindung an den Hauptsatz auf, was dazu führen kann, dass Lernende die Grenze zwischen Haupt- und Nebensatz nur schwer wahrnehmen (vgl. Feilke, 2011, S. 349). Die im DeChiLKO-Korpus häufige Fehlschreibung **das* statt *dass* trotz korrekter Kommasetzung könnte darauf hindeuten, dass die Lernenden beim Diktieren die gehörten Wörter lediglich phonetisch verschriftlicht haben, ohne die zugrunde liegende syntaktische Struktur ausreichend zu analysieren und zu verstehen.

Im Gegensatz zur häufigen Falschschreibung von <das> für <dass> treten die Fehlschreibungen von *dass* statt *das* im Korpus (siehe Beispiel 4) nur 10-mal auf.

4. [word]: [...], aber dass kann man verhindern.
 [ZH]: [...], aber das kann man verhindern.
 (DeChiLKO_Prüfung2017: D_NW_2017_D_006 (tokens 115–136))

Die umgekehrte Fehlschreibung *dass* statt *das* kann unter dem Aspekt der Phonem-Graphem-Korrespondenzen als eine regelhafte, aber grammatisch nicht kontrollierte Schreibung verstanden werden (vgl. Feilke, 2011, S. 349). Die Vokalkürze bei *das* wird, wie bei anderen einsilbigen Synsemantika (z. B. *im*, *an*, *in*, *was*), nicht durch eine Doppelkonsonanz gekennzeichnet. Die Konjunktion *dass* hingegen weist nach dem silbischen Prinzip trotz Einsilbigkeit die Doppelkonsonanz auf (vgl. Fuhrhop, 2021, S. 87). Die Lernenden könnten beim Schreiben primär die PGK-Regeln für kurze Vokale vor Konsonantenclustern beachtet haben, ohne eine hinreichende syntaktische Analyse der Funktion von *das* bzw. *dass* durchzuführen.

Bei der *man*-*Mann*-Schreibung zeigen alle 14 Fehlschreibungen im Korpus die Ersetzung des Pronomens *man* durch eine Form von *Mann* (**mann*). Dabei wird der Großteil (n = 8) der Fehlschreibungen mit kleingeschriebenem Initial als **mann* realisiert (siehe Beispiel 5). Seltener (n = 5) tritt die Ersetzung durch das großgeschriebene Substantiv *Mann* auf (siehe Beispiel 6).

5. [word]: [...], wenn mann etwas nicht gleich kann.
 [ZH]: [...], wenn man etwas nicht gleich kann.
 (DeChiLKO_Prüfung2017: A_O_2017_D_010 (tokens 33–44))
6. [word]: [...], aber das kann Mann verhindern.
 [ZH]: [...], aber das kann man verhindern.
 (DeChiLKO_Prüfung2017: A_O_2017_D_010 (tokens 113–123))

Ähnlich wie bei der fälschlichen *dass*-Schreibung lassen sich die Fehlschreibungen **mann* statt *man* unter Berücksichtigung der Phonem-Graphem-Korrespondenz als phonetisch motivierte, regelhafte Schreibungen interpretieren. Die Verdopplung des Konsonanten <nn> kann dabei als Markierung der Vokalkürze in /man/ verstanden werden. Dieses Phänomen deutet darauf hin, dass die Lernenden möglicherweise die reguläre orthographische Regel zur Markierung von Kurzvokalen durch Konsonantenverdopplung auf das Pronomen *man* übertragen haben, obwohl dies hier semantisch nicht korrekt ist.

Die Fehlschreibung mit großgeschriebenem *Mann* lässt sich darauf zurückführen, dass die Lernenden das Indefinitpronomen *man* fälschlicherweise mit dem ihnen möglicherweise vertrauten Substantiv *Mann* assoziierten. Dies ist möglicherweise auf eine mangelnde Differenzierung zwischen Pronomen und Substantiv zurückzuführen, die durch die phonologische Identität der beiden Formen im Deutschen verstärkt wird. Erschwerend kommt hinzu, dass es im L1 Chinesisch keine direkte Entsprechung für das deutsche Indefinitpronomen *man* gibt. Wenn man sich, wie im Beispielsatz „Computer sind auch geduldiger als Lehrer, wenn man etwas nicht gleich kann“, auf Menschen im Allgemeinen bezieht, wird im Chinesischen das Wort „人们“ (*rénmen*) verwendet, das wörtlich „Menschen“ oder „Personen“ bedeutet und durch das Pluralsuffix „们“ (*men*) die Allgemeinheit ausdrückt. Diese unterschiedliche Konzeptualisierung kann die korrekte Verwendung und Schreibung des deutschen *man* zusätzlich erschweren.

Zusammenfassend verdeutlichen die Fehler bei den Homophonpaaren *dass/das* und *man/Mann*, dass für die chinesischen Lernenden eine Verbesserung der tiefgehenden syntaktischen Analyse sowie des Verständnisses der Wortart und Funktion im Satz erforderlich ist. Die rein phonetische Verschriftung oder die Anwendung übergeneralisierter PGK-Regeln führt hier zu den beobachteten Fehlschreibungen.

4.6 Der sonstige Bereich

Der sonstige Bereich umfasst im DeChiLKo-Korpus im Wesentlichen zwei spezifische Annotationsebenen sowie eine Auffangkategorie:

- [Son_FW] (Fremdwortschreibung): Abweichungen bei der korrekten Schreibung von Fremdwörtern
- [Son_Gra-Ab] (Grammatisch abzuleitende Fehler): Fehler, die direkt auf eine fehlerhafte Anwendung grammatischer Regeln zurückzuführen sind und sich in der Schreibung niederschlagen (z. B. falsche Flexionsendungen, die zu einer orthographisch inkorrekten Wortform führen, auch wenn die Stammform korrekt sein mag)²⁸

²⁸ Es ist an dieser Stelle hervorzuheben, dass die Lernertexte im DeChiLKo-Korpus auf Diktataufgaben basieren. Dies stellt eine besondere Herausforderung für die eindeutige Klassifizierung bestimmter Fehler dar. So ist es beispielsweise bei einer Fehlschreibung wie **in Haus* statt *im Haus* oft schwierig, allein anhand der Oberflächenform zu entscheiden, ob es sich primär um eine phonographische Abweichung oder um einen Deklinationsfehler handelt. Um dieser Komplexität gerecht zu werden und eine umfassendere Analyse zu ermöglichen, wurde im Annotationsprozess entschieden, solche Fälle mehrfach zu annotieren: Abweichungen, die potenziell sowohl eine orthografische als auch eine grammatische Ursache haben könnten, werden sowohl auf den entsprechenden spezifischen Annotationsebenen (z. B. [Graph_PGK]) als auch auf der Ebene [Son_Gra-Ab] markiert. Diese Vorgehensweise erlaubt eine differenziertere Untersuchung der Fehlerursachen und der Interaktion verschiedener sprachlicher Ebenen.

- [SON] (Sonstige Fehler): Abweichungen, die keine klare Verbindung zur Zielhypothese aufweisen oder anderweitig nicht klassifizierbar sind. Beispiele hierfür sind sinnentstellende Wortverwechslungen (etwa *<die> für *ist*, *<dass> für *Leute*), unvollständige Wörter (z. B. Lernende schreiben lediglich die ersten Buchstaben eines Wortes wie *<lo> für *Leute* oder *<g> für *Garten*) oder die fälschliche Verwendung von Wörtern aus anderen Sprachen (z. B. *<Bath> für *Bad*).

4.6.1 Fehler bei der Fremdwortschreibung [Son_FW]

Die deutsche Sprache hat im Laufe ihrer Geschichte zahlreiche Wörter aus anderen Sprachen entlehnt. Während im 19. Jahrhundert das Französische die häufigste Quelle für Fremdwörter war, werden heutzutage die meisten Fremdwörter aus dem Englischen übernommen (vgl. Fuhrhop, 2021, S. 33). Diese Entwicklung spiegelt sich auch in den im DeChiLKo-Korpus vorkommenden Fremdwörtern wider, die ebenfalls hauptsächlich aus dem Englischen stammen (z. B. *Hobby*, *Computer* usw.). Vereinzelt finden sich auch bekannte, aus dem Französischen entlehnte Fremdwörter in den Zielhypothesen (z. B. *Café*, *Toilette*).

Insgesamt wurden im DeChiLKo-Korpus lediglich 21 Abweichungen bei der Fremdwortschreibung in 21 Lernertexten festgestellt. Dies deutet auf eine vergleichsweise geringe Fehleranfälligkeit in diesem spezifischen Bereich hin, was jedoch auch durch die Auswahl der Diktattexte bedingt sein kann, die nur eine begrenzte Anzahl an Fremdwörtern enthielten.

Die häufigsten Fehler ($n = 7$) betreffen das Wort *Hobby* in seiner Pluralform *Hobbys*. Statt der korrekten deutschen Pluralform *Hobbys* greifen die Lernenden hier häufig auf die englische Pluralform *hobbies* zurück. Interessanterweise wird diese englische Pluralform im Korpus einmal als *<hobbies> (kleingeschrieben) realisiert, während sie sonst unter Beibehaltung der englischen Pluralendung, aber mit deutscher Großschreibung am Wortanfang, als *<Hobbies> erscheint.

Die Übernahme der englischen Pluralbildung bei *Hobby* zu *<hobbies> deutet auf einen klaren Transfereffekt aus der L2 Englisch hin, der im multikompetenten Repertoire der Lernenden verankert ist. Die gleichzeitige Anwendung der deutschen Regel zur Substantivgroßschreibung in der Form *<Hobbies> ist ein interessantes Beispiel für eine hybride Schreibstrategie. Hier kombinieren die Lernenden orthographische Merkmale aus zwei ihnen bekannten Sprachsystemen (englische Pluralmorphologie und deutsche Großschreibungsregel), um eine plausible, wenn auch nicht normgerechte, Schreibweise für das deutsche Wort zu generieren. Dies unterstreicht, wie Lernende aktiv ihr gesamtes sprachliches Wissen und ihr verfügbares orthographisches Repertoire beim Schreiben nutzen.

Im Vergleich zur relativ einheitlichen und systematischen Abweichung bei der Schreibweise von *<Hobbies> stellt das aus dem Französischen entlehnte Wort *Toilette* (/toa'letə/) eine deutlich größere Herausforderung für die Lernenden dar. Dieses Fremdwort wird in den Lernertexten in vielfältigen Varianten verschriftlicht, beispielsweise als *<Tualäter>, *<Turlete>, *<Turleiter>, *<Tualetter> oder *<Turlette>. Besonders auffällig sind dabei Fehlschreibungen, die auf eine fehlerhafte Perzeption und/

oder graphemische Umsetzung von Laut-Buchstaben-Verbindungen zurückzuführen sind, die im Deutschen so nicht existieren oder anders realisiert werden. Für den im Deutschen ungewöhnlichen Diphthong /oa/ nehmen die Lernenden möglicherweise den ähnlich klingenden, im hinteren Vokalbereich angesiedelten deutschen Vokal /u/ oder /u:/ wahr und verschriftlichen ihn als <u> (z. B. in *<Turlete>). Der Vokal /a/ in der ersten Silbe wird hingegen entweder gemäß der deutschen Phonem-Graphem-Korrespondenz als <a> wiedergegeben oder – möglicherweise in Analogie zu im Deutschen häufig vorkommenden zentralisierenden Diphthongen – als Vokal mit einem nachfolgenden <r> verschriftlicht (z. B. in *<Turleiter>). Diese Übertragung deutscher Schreibmuster auf das französische Fremdwort belegt eine Tendenz der Lernenden, fremdsprachliche Laut-Buchstaben-Beziehungen an die ihnen vertrauten Muster der Orthographie anzupassen. Ähnliche Übergeneralisierungen der deutschen Phonem-Graphem-Korrespondenz treten auch bei der Schreibung anderer Fremdwörter wie *Café* und *Computer* auf. So wird das französische *Café* häufig als *Kaffee* und das aus dem Englischen entlehnte Wort *Computer* als *<Komputer> verschriftlicht, was auf eine Anpassung an die deutschen Aussprache- und Schreibgewohnheiten hindeutet.

Darüber hinaus weisen die Fehlschreibungen wie *<Turlete> oder *<Tualäter> darauf hin, dass chinesische Lernende generelle Schwierigkeiten mit der Konsonantenverdopplung im Deutschen haben, wobei hier die Verdopplung des Konsonanten <tt> im Silbengelenk von *Toilette* häufig einfach geschrieben wird.

Interessanterweise treten auch zwei Fehlschreibungen bei dem aus dem Chinesischen entlehnten Fremdwort *Tai-Chi* auf. Dieses chinesische Schattenboxen wird in der Originalsprache als 太极 (/tʰaɿ˥˩ tei/) bezeichnet und in der *Pinyin*-Transkription als *Taiji* wiedergegeben. Im Deutschen wird es jedoch orthographisch und phonetisch angepasst, wobei der stimmlose, nicht aspirierte Palatal /tɕ/ (*Pinyin* <j>) an das deutsche Lautsystem angenähert wird. Dabei erfolgt die Integration des chinesischen Konsonanten /tɕ/ als /ɕ/, der im Deutschen durch die Grapheme <ch> verschriftlicht wird. Im Korpus sind jedoch die Fehlschreibungen *<Taiji> und *<Taiqi> zu finden. Die Schreibweise *<Taiji> resultiert vermutlich aus der direkten Übernahme der *Pinyin*-Form. Die Lernenden haben das Wort *Tai-Chi* beim Hören korrekt verstanden und eine Verbindung zu dem chinesischen Wort 太极 (*taiji*) hergestellt. Infolgedessen haben sie das Wort entsprechend der *Pinyin*-Transkription verschriftlicht, ohne die orthographische Anpassung des Fremdwortes zu berücksichtigen. Dies ist ein klarer Fall von L1-Transfer (*Pinyin*). Die Schreibweise *<Taiqi> hingegen deutet darauf hin, dass der/die Lerner/in das Wort *Tai-Chi* wahrscheinlich nicht vollständig mit dem chinesischen Ursprungswort identifiziert hat und es ausschließlich auf Basis der lautlichen Wahrnehmung verschriftlichte. Dabei wurde der deutsche Laut [ɕ] möglicherweise mit der aspirierten chinesischen Lautverbindung [tɕʰ] (*Pinyin*: <q>) verwechselt, was zu einer Übertragung nach der Phonem-(*Pinyin*)Graphem-Korrespondenz <q> für /tɕʰ i/ führte.

Die Analyse der Abweichungen bei der Fremdwortschreibung bekundet, dass die Integration von Lehnwörtern aus unterschiedlichen Sprachquellen für chinesische Deutschlernende eine besondere Herausforderung darstellt. Während Wörter aus dem Englischen wie *Hobby* und *Computer* durch Interferenzen und hybride Schreibstra-

tegien geprägt sind, führen französische Lehnwörter wie *Toilette* und *Café* aufgrund ihrer spezifischen Lautstruktur zu vielfältigen, oft an deutschen Mustern angepassten Schreibweisen. Die Fehlschreibungen des aus dem Chinesischen stammenden Fremdwortes *Tai-Chi* verdeutlichen zudem, wie die L1 und das *Pinyin*-System das Schreibverhalten auch bei der Verschriftung von Wörtern mit bekanntem Ursprung beeinflussen können.

Diese Ergebnisse unterstreichen, dass die Lernenden bei der Fremdwortschreibung auf ihr gesamtes multikompetentes Repertoire zurückgreifen und ihr orthographisches Verhalten sowohl von den Regeln der deutschen Orthographie als auch von ihrem bereits vorhandenen sprachlichen Wissen aus L1 und L2 geprägt ist. Die dabei entstandenen, oft variablen und kontextabhängigen Schreibweisen (z. B. *<Tualäter>, *<Turlete>, *<Turleiter>, *<Tualetter> oder *<Turlette> für *Toilette*; *<Taiji>, *<Taiqi> für *Tai-Chi*) implizieren somit, dass Transfereffekte keine statischen Prozesse sind, sondern sich dynamisch aus dem Zusammenspiel der verschiedenen sprachlichen Systeme und den spezifischen Anforderungen der Schreibaufgabe neu strukturieren. (Vgl. Larsen-Freeman, 2017, S. 26).

4.6.2 Grammatisch abzuleitende Fehler [Son_Gra]

Im Einklang mit dem Prinzip der mehrschichtigen Annotation der vorliegenden Untersuchung dient die Annotationsebene [Son_Gra] einer zusätzlichen Analyse dazu, ob Fehlschreibungen grammatisch abzuleiten sind. In einem Diktattext ist es oft schwierig zu bestimmen, ob eine Abweichung wie **im See* statt *in See* lediglich auf einer Verwechslung der Phoneme /m/ und /n/ beruht oder ob sie auch als Deklinationsfehler interpretiert werden soll. Streng genommen handelt es sich bei Deklinationsfehlern um ein morphologisches bzw. morphosyntaktisches Phänomen. Da solche Fehler jedoch oft aus einer falschen Analyse der syntaktischen Beziehungen im Satz resultieren, werden sie im Rahmen dieses Kapitels mitbetrachtet. Aus diesem Grund werden auf dieser Annotationsebene Abweichungen nicht nur unter dem Gesichtspunkt der Phonem-Graphem-Korrespondenz betrachtet, sondern auch aus grammatischer Perspektive kategorisiert. Je nach Fehlerart werden die Abweichungen als Deklinationsfehler (Dek), Konjugationsfehler (Kon) oder Fehler bei der Auswahl von Präpositionen (PP) unterschieden. Deklinationsfehler beziehen sich auf die Beugung von Nomina, Pronomen, Adjektiven oder Artikelwörtern. Sie treten auf, wenn die Wortform nicht korrekt in Bezug auf Kasus, Numerus oder Genus dekliniert wird. Konjugationsfehler betreffen die Flexionsformen von Verben. Diese Fehler entstehen, wenn die Verbform nicht korrekt nach Person, Numerus, Tempus oder Modus realisiert wird (vgl. Eisenberg, 2016, S. 132–133).

Insgesamt sind 822 grammatisch abzuleitende Abweichungen in 267 Lernertexten im gesamten Korpus zu finden, dabei beträgt der Fehlerwert pro 1.000 Tokens 25,95.

Tabelle 40: Überblick über grammatisch abzuleitende Abweichungen

Grammatisch abzuleitende Abweichungen insgesamt	Summe aller Tokens	Grammatisch abzuleitende Abweichungen pro 1.000 Tokens
822	31.674	25,95

Die folgende Abbildung stellt die Fehlerverteilung nach Fehlerarten dar.

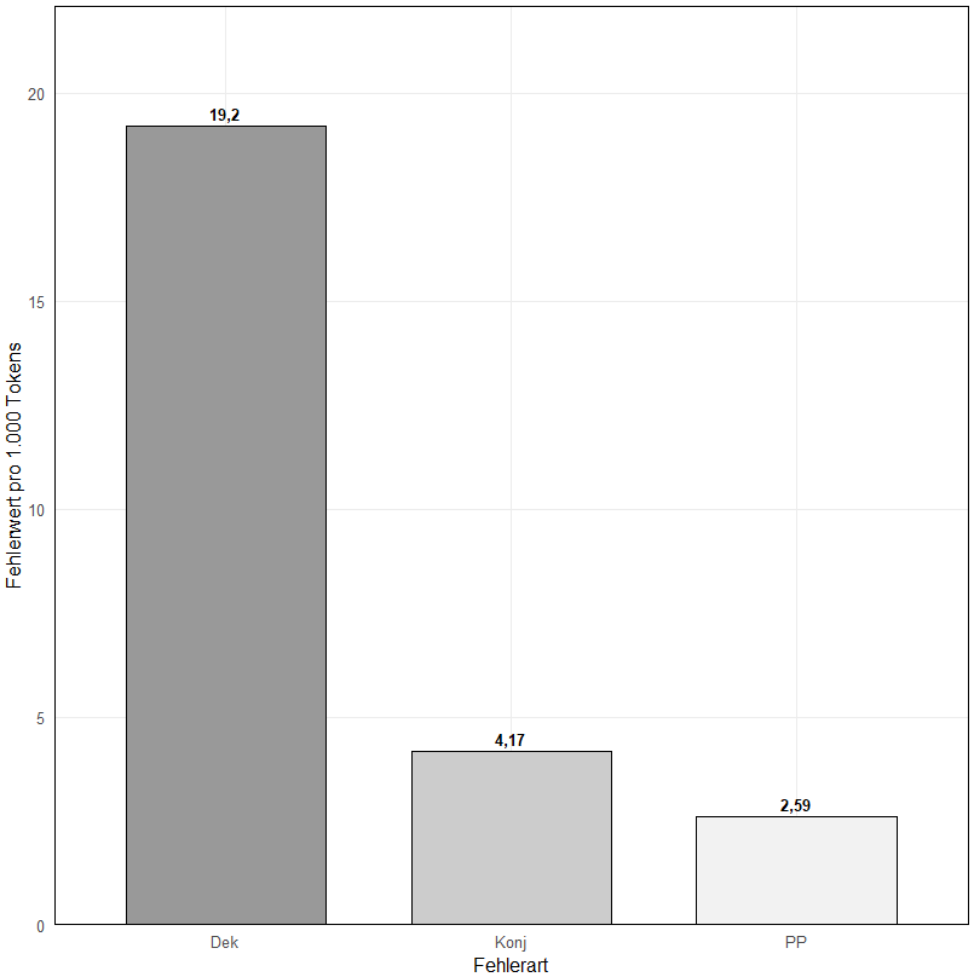


Abbildung 32: Grammatisch abzuleitende Fehler nach Fehlerarten; Fehlerwert pro 1.000 Tokens

Deutlich kommen Deklinationsfehler mit einem Fehlerwert von 19,20 pro 1.000 Tokens am häufigsten im Korpus vor. Mit großem Abstand stehen Konjugationsfehler auf dem zweiten Platz, der Fehlerwert beträgt 4,17. Fehler aufgrund der Auswahl von Präpositionen treten im Korpus selten auf. Das Ergebnis, dass Deklinationsfehler eine

wesentliche Fehlerquelle in der Lernersprache darstellen, deckt sich mit zahlreichen Befunden der DaF-Forschung (vgl. u. a. Diehl, 2000; Grieshaber, 2016). Lay (2017, S. 35) betont insbesondere die Lernschwierigkeiten bei der Deklination und Konjugation im Deutschen für chinesische Deutschlernende, da ihre Erstsprache Chinesisch eine weitgehend isolierende Sprache ist, in der es keine Flexion im Sprachsystem gibt.

Wie es bereits im Kapitel 4.2 dargelegt wurde, finden sich im gesamten Korpus insgesamt 4.529 Abweichungen in der Phonem-Graphem-Korrespondenz (Annotationsebene [Graph_PGK]). Die Analyse der Annotationsebene [Son_Gra-Ab] zeigt zudem, dass 979 dieser Abweichungen aus einer grammatischen Perspektive erklärbar sind. Die folgende Tabelle stellt die Anzahl der Fehler verschiedener Fehlertypen auf der Ebene der Phonem-Graphem-Korrespondenz dar und zeigt zugleich, wie viele dieser Fehler auch als grammatisch bedingte Abweichungen eingeordnet werden können.

Tabelle 41: Übersicht der Fehlerverteilung auf der Annotationsebene [Graph_PGK] und deren grammatische Klassifikation mit prozentualen Anteil

[PGK]	Gesamtzahl der PGK-Abweichungen	Anzahl der Deklinationsfehler (Dek)	Anzahl der Konjugationsfehler (Konj)	Anzahl der fehlerhaften Auswahl der Präpositionen (PP)	Gesamtzahl der grammatisch abzuleitenden Abweichungen	prozentualer Anteil
K-	785	145	37	2	184	23 %
*K	727	133	11	29	173	24 %
*V	726	39	10	23	72	10 %
K+	573	101	55	4	160	28 %
RV-	483	154	46	0	200	41 %
*RV	483	87	0	0	87	18 %
RV+	283	23	21	2	46	16 %
*VR	154	4	0	35	39	25 %
V-	93	0	0	0	0	0 %
VR-	88	12	0	2	14	16 %
VR+	64	1	0	1	2	3 %
V+	37	0	0	0	0	0 %
GV	33	2	0	0	2	6 %
Gesamtzahl	4.529	701	180	98	979	22 %

Wenn man die Anteile der grammatisch abzuleitenden Fehler in Relation zu den Abweichungen der Graphem-Phonem-Korrespondenz betrachtet, fällt die Fehlerart „fehlende Reduktionsvokalgraphem(e)“ (RV-) besonders ins Auge. Mit einem Anteil von

41 % grammatisch erklärbarer Abweichungen weist diese Kategorie die höchste relative Häufigkeit grammatisch bedingter Abweichungen auf. Eine detaillierte Analyse macht deutlich, dass diese Fehler hauptsächlich auf das Fehlen von Deklinationseendungen zurückzuführen sind, insbesondere bei Adjektiven und Substantiven, wie es in den folgenden Beispielen sichtbar wird:

7. [word]: Durch verschieden Arten von Belohnung werden die Kinder [...]
 [ZH]: Durch verschiedene Arten von Belohnungen werden die Kinder [...].
 (DeChiLKo_Prüfung2017 > C_N_2017_D_006 (tokens 1–13))
8. [word]: Sie brauchen Geschenk für sie.
 [ZH]: Sie brauchen Geschenke für sie.
 (DeChiLKo_Erwerb > LWY20211227 (tokens 6–27))

Diese Beispiele demonstrieren, dass Lernende häufig Schwierigkeiten mit der korrekten Markierung von Kasus und Numerus haben, insbesondere bei der Pluralbildung und der Anpassung von Adjektivendungen an das Substantiv.

Außerdem ist der Fehlertyp „überflüssige Konsonantengraphem(e)“ (K+) mit einem Anteil von 28 % grammatisch erklärbarer Abweichungen ebenfalls bemerkenswert. Die Analyse unter Berücksichtigung der Zielhypothese [ZH] zeigt, dass die meisten Abweichungen auf das Einfügen überflüssiger Deklinationseendungen zurückzuführen sind, wie es die folgenden Beispiele veranschaulichen:

9. [word]: Natürlich gibt es auch Problemen.
 [ZH]: Natürlich gibt es auch Probleme.
 (DeChiLKo_Erwerb > LWY20220623 (tokens 96–115))
10. [word]: Durch verschieden~~en~~ alten Vorbelohnungen werden die Kinder [...].
 [ZH]: Durch verschiedene Arten von Belohnungen werden die Kinder [...].
 (DeChiLKo_Prüfung2017 > D_NW_2017_D_009 (tokens 1–13))

Darüber hinaus treten viele überflüssige Konsonantengrapheme auch im Zusammenhang mit der falschen Konjugation von Verben auf:

11. [word]: Aus diesen grüden sollten man den Computer nicht von Anfang an Ablemen.
 [ZH]: Aus diesen Gründen sollte man den Computer nicht von Anfang an ablehnen.
 (DeChiLKo_Prüfung2017 > B_N_2017_D_003 (tokens 55–76))

In einer Korpusstudie zu orthographischen Fehlern in den Abiturtextrn stellen Berg, Hartmann und Claeser (2023, S. 12) fest, dass Flexionssuffixe eine signifikant höhere Wahrscheinlichkeit für Auslassungen von Endbuchstaben aufweisen als derivationale Suffixe. Diese Beobachtung deckt sich mit den Ergebnissen des vorliegenden Korpus, in dem 19,3 % aller Abweichungen der Phonem-Graphem-Korrespondenz auf Flexionsfehler (einschließlich Konjugations- und Deklinationsfehler) zurückzuführen sind, während nur 6,8 % der Abweichungen auf Derivationssuffixe (Annotationsebene [Morph_AffixS]) entfallen.

Obwohl über 20 % aller Abweichungen der Phonem-Graphem-Korrespondenz in den lernersprachlichen Diktattexten als grammatisch bedingte Fehler, insbesondere als Deklinationsfehler, einzuordnen sind, besteht der Großteil der Abweichungen weiterhin aus Verstößen gegen die phonographische Regularität der deutschen Sprache, wie bereits in Kapitel 4.2 ausführlich analysiert. Diese Verstöße verdeutlichen, dass die korrekte Umsetzung der Laut-Buchstaben-Zuordnung ebenso eine Herausforderung für chinesische Deutschlernende darstellt.

4.6.3 Sonstige Fehler [SON]

Auf der sonstigen Annotationsebene [SON] werden jene Abweichungen in den Lerner-texten erfasst, die keiner anderen spezifischen orthographischen Kategorie eindeutig zugeordnet werden können. Im gesamten DeChiLKo-Korpus wurden 222 Fehler dieser Kategorie in 118 Lernertexten identifiziert, was einem Fehlerwert von 7,01 pro 1.000 Tokens entspricht.

Der überwiegende Großteil dieser Abweichungen ($n = 211$) wurde mit dem allgemeinen Tag *SON* markiert. Dies bedeutet, dass die verschriftete Worteinheit entweder keine erkennbare semantische oder graphemische Verbindung zur Zielhypothese aufweist (sogenannte sinnveränderte Fehler) oder unvollständig ist und beispielsweise nur mit den Anfangsbuchstaben notiert wurde (z. B. *<le> für *Leute* oder *<g> für *Garten*). Darüber hinaus wurden 11 Fehler (*Eng*) festgestellt, bei denen englische Wörter fehlerhaft anstelle der deutschen Zielwörter verwendet wurden (vgl. Tabelle 42).

Tabelle 42: Überblick über sonstige Fehler im gesamten Korpus

Abweichungen	Anzahl	Summe aller Tokens	Abweichungen pro 1.000 Tokens
SON	211	31.674	6,66
Eng	11		0,35

Die vermehrt auftretenden unvollständigen Wörter können auf die Anwendung von „Tipps für das Diktat“ zurückgeführt werden. In der Prüfungsvorbereitung wird den Lernenden häufig empfohlen, Wörter, die als bekannt vorausgesetzt werden, zur Zeitersparnis während des Diktats zunächst nur mit den Anfangsbuchstaben zu notieren und diese erst später zu vervollständigen (vgl. Xu et al., 2016, S. 59). Unter den spezifischen Prüfungsbedingungen des Diktats kann es jedoch passieren, dass Lernende aufgrund von Zeitdruck oder Stress vergessen, diese fragmentarisch notierten Wörter im Nachhinein zu ergänzen.

Auch bei den sinnverändernden Abweichungen spielt der Leistungsdruck in der Prüfungssituation eine entscheidende Rolle. Wenn es den Lernenden während des Diktats aufgrund von Zeitmangel oder Verständnisschwierigkeiten nicht gelang, das gehörte Wort korrekt zu verschriften, griffen sie häufig auf ein ihnen bekanntes Wort zurück. Diese Strategie diene möglicherweise dazu, die entstandene Lücke zu füllen und zumindest teilweise eine schriftsprachliche Leistung zu erbringen. Besonders häu-

fig wurde dabei ein Wort gewählt, das denselben Initialbuchstaben wie die intendierte Zielhypothese aufweist. Ein Beispiel hierfür ist Beispiel 12, in dem anstelle von *Spiel* das Wort *Stellen* geschrieben wurde:

12. [word]: [...]. In dem man die ***Stellen** und *Lernen *dauern vorher genau *fesslegt.
 [ZH]: [...], indem man die **Spiel**- und Lerdauer vorher genau festlegt.
 (DeChiLKo_Prüfung2017 > D_NW_2017_D_003 (tokens 108–118))

Darüber hinaus sind Fehlschreibungen bemerkenswert, die stark auf Transfereffekte aus der ersten Fremdsprache Englisch hindeuten (Kategorie [Eng]), auch wenn derartige Fehlererscheinungen im Korpus insgesamt nur selten auftreten (Fehlerwert 0,35/1.000 Tokens).

13. [word]: Immer mehr Deutsche gehe ins ***Swimbat**.
 [ZH]: Immer mehr Deutsche gehen ins **Schwimmbad**.
 (DeChiLKo_Prüfung2019 > B_N_2019_D_003 (tokens 87–101))
14. [word]: ***freeziet** *Müsst immer *Kosten.
 [ZH]: **Freizeitvergnügen** muss nicht immer Geld kosten.
 (DeChiLKo_Prüfung2019 > C_NW_2019_D_008 (tokens 164–178))

Im Beispiel 13 wird deutlich, wie Transfereffekte aus L2 wirksam werden. Das Erstglied des Kompositums *Schwimmbad* scheint hier durch das englische Wort *swim* ersetzt worden zu sein. Entsprechend der deutschen Wortbildungslogik greift der/die Lerner/in auf den reinen Verbstamm *swim* zurück, anstatt die im Englischen für solche Phrasen übliche Form *swimming* (wie in *swimming pool*) zu nutzen. Darüber hinaus wurde das hybride Kompositum *<Swimbat> gemäß den deutschen orthographischen Regeln für Substantive großgeschrieben, und die beiden Wortglieder wurden zusammengeschrieben. Die Fehlerschreibung *{bat} anstelle von {bad} im Zweitglied verdeutlicht zudem, dass der/die Lernende noch Schwierigkeiten mit der korrekten Schreibung der deutschen Auslautverhärtung hat. Diese Fehlschreibung ist ein gutes Beispiel für die Interaktion verschiedener sprachlicher Systeme im multikompetenten Repertoire der Lernenden. Hier wird englisches Lexikon mit deutschen Wortbildungs- und Orthographieregeln kombiniert, was zu einer emergenten, individuell geprägten lernersprachlichen Orthographie führt.

Solche fälschlich angewandten englischen Wörter treten im DeChiLKo-Korpus insbesondere dann auf, wenn sie in Form und in Bedeutung Ähnlichkeiten mit der Zielhypothese in der deutschen Sprache aufweisen, wie etwa *swim* und *Schwimm*. Im Beispiel 14 wurde das englische Adjektiv *free* als Erstglied im Kompositum *Freizeitvergnügen* verwendet. Da dieses komplexe deutsche Wort für die Lernerin oder den Lerner vermutlich ein fremdes und wenig vertrautes Lexem darstellte, wurde zudem das Endglied {vergnügen} ausgelassen. Die weiteren Fehlschreibungen *<Müsst> und *<Kosten> im selben Satz deuten zusätzlich darauf hin, dass die Lernerin oder der Lerner die Regeln der Groß- und Kleinschreibung der deutschen Sprache noch nicht vollständig verinnerlicht hat.

Zusammenfassend verdeutlicht die Analyse der sonstigen Abweichungen, dass die meisten dieser Fehler auf unvollständig geschriebene Wörter oder sinnverändernde Abweichungen zurückzuführen sind. Diese können maßgeblich durch Prüfungsdruck und Zeitmangel beeinflusst werden. Dies deutet darauf hin, dass Lernende in Stresssituationen Strategien wie die Verwendung von Anfangsbuchstaben oder formal bzw. klanglich ähnlichen, aber semantisch abweichenden, vertrauten Wörtern anwenden, um Unsicherheiten auszugleichen. Die wenigen, jedoch bemerkenswerten Fälle der fehlerhaften Verwendung englischer Elemente weisen darauf hin, dass sprachübergreifende Interferenzen zwischen Sprachen vor allem dann auftreten, wenn sich die Form und Bedeutung der Wörter ähneln. Die Lernenden greifen beim Schreiben auf ihr gesamtes sprachliches Repertoire zurück. Dies spiegelt sich in den hybriden und oft individuell geprägten Schreibungen deutlich wider.

5 Relation der Orthographieleistung zu verschiedenen Faktoren

Um die Faktoren wie Lerndauer, Erhebungsbedingungen, Hochschulkategorie und -typ sowie Erhebungsregion inferenzstatistisch zu überprüfen und zu kontrollieren, werden generalisierte lineare gemischte Modelle (GLMM) und lineare gemischte Modelle (LMM) mithilfe des R-Pakets *lme4* (Bates et al., 2015) berechnet. Da die Lernertexte im Erwerbskorpus und im Prüfungskorpus durch unterschiedliche Erhebungsverfahren gesammelt wurden, werden die Einflussfaktoren für beide Subkorpora separat betrachtet.

5.1 Faktoren im Prüfungskorpus

Die Lernertexte im Prüfungskorpus stammen aus den landesweiten PGG-Prüfungen der Jahre 2017 und 2019. Die Diktataufgaben unterschieden sich in beiden Erhebungsjahren sowohl hinsichtlich des verwendeten Wortschatzes als auch der Textlänge. Gemäß dem Curriculum (Bildungsministerium, 2006) des Bachelorstudiengangs Deutsch sollen alle Studentinnen und Studenten der Germanistik oder verwandter Fächer am Ende ihres zweiten Studienjahres diese Prüfung ablegen. Studierende, die die Prüfung zu diesem Zeitpunkt nicht bestehen, haben die Möglichkeit, sie in den folgenden Jahren zu wiederholen, sofern sie weiterhin an der Universität immatrikuliert sind. Das Prüfungskorpus umfasst insgesamt 195 Lernertexte. Die größte Gruppe bilden dabei die 141 Texte von Studierenden des zweiten Studienjahres, die die Prüfung turnusgemäß ablegten (reguläre Teilnehmende). Die verbleibenden 54 Texte stammen von Studierenden des dritten ($n = 40$) sowie des fünften und sechsten Studienjahres ($n = 14$), die die Prüfung wiederholten (Wiederholer*in). Aufgrund dieser Struktur wurde für die statistische Analyse nicht mehr strikt nach Studienjahren differenziert, sondern der Faktor „Prüfungsstatus“ als binäre Variable operationalisiert: Er unterscheidet zwischen regulären Teilnehmer*innen (2. Studienjahr) und Wiederholer*innen (3. Studienjahr und höher).

Zudem wurden die teilnehmenden Hochschulen bereits bei der Datenerfassung in fünf Kategorien (A bis E) eingeteilt (siehe Kapitel 3.2.1.2). Diese Kategorisierung reflektiert das Renommee der Hochschulen, das eng mit der für die Zulassung erforderlichen Punktzahl bei der nationalen Hochschulaufnahmeprüfung *Gaokao* korreliert (vgl. Goldberger, 2017). Setzt man die in der Regel maximal erreichbare Gesamtpunktzahl von 750 Punkten als Maßstab an, wird die Dimension dieser Leistungsunterschiede deutlich. So benötigte beispielsweise ein Studierender aus der Provinz Shandong im Jahr 2023 für das Germanistikstudium an der Nanjing Universität (Kategorie A) mindestens 643 Punkte, während für dasselbe Fach an der Shandong Vocational University of Foreign Affairs (Kategorie E) 277 Punkte ausreichten. Es ist daher davon auszugehen, dass die Hochschulkategorien mit der allgemeinen Leistungsfähigkeit der Studierenden in Zusammenhang stehen.

Darüber hinaus spielte der Standort der Institutionen eine wichtige Rolle. Um möglichst repräsentative Daten zu gewährleisten, wurden Hochschulen aus sechs geografischen Regionen Chinas ausgewählt (Nord-, Nordwest-, Nordost-, Ost-, Südwest-, Zentral- und Südchina). Da sich Infrastruktur, Lehrressourcen und Unterrichtsmethoden zwischen diesen Regionen unterscheiden können, wurde der Faktor Region ebenfalls in die Modellierung einbezogen.

5.1.1 Statistische Analyse mittels Hurdle-Modells

Zur Analyse des durchschnittlichen Fehlerwertes pro 1.000 Tokens im Prüfungskorpus wurde ein zweistufiges Modellierungsverfahren gewählt: das Hurdle-Modell (Cameron & Trivedi, 2013; Zeileis, Kleiber & Jackman, 2008). Diese methodische Entscheidung gründet auf der spezifischen Verteilung der abhängigen Variable: Ein beträchtlicher Anteil von 31,7% der Beobachtungen weist einen Fehlerwert von Null auf. Dies bedeutet, dass in diesen Texten keine Fehler der entsprechenden Kategorie auftraten. Das Hurdle-Modell erlaubt eine adäquate Berücksichtigung dieser strukturellen Nullwerte, indem es den Prozess des Fehlerauftretens und den Prozess der Fehlerausprägung (Magnitude) getrennt voneinander modelliert (vgl. Feng, 2021).

Das Hurdle-Modell zerlegt die Analyse in zwei separate Stufen: In der ersten Stufe wird die Wahrscheinlichkeit modelliert, ob in einem Lernertext überhaupt ein orthographischer Fehler auftritt (Fehlerwert > 0), im Vergleich zum Ausbleiben eines Fehlers (Fehlerwert = 0). Hierfür kommt ein generalisiertes lineares gemischtes Modell (GLMM) mit binomialer Fehlerverteilung und einer Logit-Verknüpfungsfunktion zum Einsatz. Die zweite Stufe analysiert die Höhe des Fehlerwertes ausschließlich für jene Fälle, in denen ein orthographischer Fehler aufgetreten ist (d. h. für alle Beobachtungen mit Fehlerwert > 0). Für diese positive Teilmenge der Daten wird der durchschnittliche Fehlerwert log-transformiert, um die Verteilungsannahmen linearer gemischter Modelle (LMM) besser zu erfüllen, und anschließend mit einem LMM analysiert.

Im Prüfungskorpus wurden die Faktoren Diktattext (kategorial: PGG 2017 vs. 2019), Prüfungsstatus (Wiederholer*innen vs. Reguläre Teilnehmer*innen), die Hochschulkategorie (A bis E) und Region als Kontrollfaktoren in beiden Stufen des Hurdle-Modells berücksichtigt. Die 18 Fehlerkategorien (Annotationsebene²⁹ [Graph_PGK], [Graph_BF], [Graph_VQM], [Graph_S-Schreibung], [Graph_SG], [Silben_AR], [Silben_SK], [Silben_ER], [Morph_Konstanz], [Morph_Affix], [Morph_Kompo], [Syn_GKS], [Syn_GZS], [Syn_dass-das], [Syn_Mann-man], [Son_FW], [Son_Gra], [SON]) wurden konsistent als zufällige Effekte (*random effects*) modelliert.³⁰ Die Entscheidung, Fehlerkategorien als zufällige Effekte zu behandeln, basiert auf mehreren methodischen Überle-

29 Um im nächsten Schritt die Interaktionsanalyse zwischen Fehlerkategorie und weiteren Faktoren durchzuführen, wurden unter den silbischen Abweichungen unter anderem das fehlende Dehnungs-*h* im Silbenkern sowie die fehlerhafte Silbengelenkschreibung als spezifische Fehlertypen berücksichtigt, während andere Fehlertypen, wie etwa falsche Anfangsränder (Tag: **ein*), nicht in die inferenzstatistische Analyse einbezogen wurden. Ähnlich wurde die Annotationsebene [Morphem], die eine Unterscheidung zwischen lexikalischen und grammatischen Morphemen vornimmt, nicht als eigenständige Fehlerkategorie in der morphologischen Schreibung bei der Berechnung berücksichtigt.

30 Die Log-Transformation des Fehlerwerts in der zweiten Stufe des Hurdle-Modells (ausschließlich für positive Fehlerwerte) dient der Annäherung an die Verteilungsannahmen linearer gemischter Modelle, wie z. B. der Symmetrisierung der Verteilung der Residuen und der Homogenisierung der Varianzen (Minimierung heteroskedastischer Effekte).

gungen. Erstens sind die Fehlerkategorien nicht unabhängig voneinander, sondern weisen Überlappungen auf. Zweitens sind die Fehlertypen im Datensatz unterschiedlich häufig vertreten. Die Modellierung als zufällige Effekte trägt diesem Umstand Rechnung, indem sie die Varianz innerhalb der Fehlerkategorien als eine zusätzliche Quelle zufälliger Variation berücksichtigt.

Durch diese zweistufige Vorgehensweise wird eine differenzierte Betrachtung der Einflussfaktoren ermöglicht. Es kann somit separat untersucht werden, welche Faktoren die Wahrscheinlichkeit des Auftretens eines orthographischen Fehlers beeinflussen und welche Faktoren dessen Ausprägung (Magnitude) bestimmen. Die Signifikanz der einzelnen Faktoren wurde mittels Likelihood-Quotienten-Tests ermittelt, indem das finale Modell in beiden Stufen jeweils mit einem Modell ohne den betreffenden Faktor verglichen wurde (vgl. Winter, 2020, S. 260–261).

Die folgende Tabelle 43 umfasst die Ergebnisse der Likelihood-Quotienten-Tests für die erste Stufe des Hurdle-Modells (Fehlerwahrscheinlichkeit).

Tabelle 43: Ergebnisse der Likelihood-Quotienten-Tests in der ersten Stufe des Hurdle-Modells

	p-Werte	Chi-Quadrat (Freiheitsgrade)	Signifikanz
Diktattext	0,0036 (<0,05)	$\chi^2 (1) = 8,446$	signifikant
Hochschulkategorie	<2,2e-16 (<0,001)	$\chi^2 (4) = 108,68$	hochgradig signifikant
Region	1,78e-06 (<0,001)	$\chi^2 (5) = 34,634$	hochgradig signifikant
Prüfungsstatus	0,0007 (<0,001)	$\chi^2 (1) = 11,41$	hochgradig signifikant

Die Ergebnisse der Likelihood-Ratio-Tests für die erste Stufe des Hurdle-Modells verdeutlichen, dass sämtliche untersuchten Faktoren einen statistisch signifikanten Einfluss auf die Fehlerwahrscheinlichkeit ausüben.

Im zweiten Teil des Hurdle-Modells (Fehlermagnitude) wurde untersucht, welche Faktoren einen signifikanten Einfluss auf die **Höhe des Fehlerwerts** haben, *sofern ein Fehler aufgetreten ist*. Die folgende Tabelle 44 fasst die Ergebnisse zusammen.

Tabelle 44: Ergebnisse der Likelihood-Quotienten-Tests in der zweiten Stufe des Hurdle-Modells

	p-Werte	Chi-Quadrat (Freiheitsgrade)	Signifikanz
Diktattext	0,0003 (<0,001)	$\chi^2 (1) = 12,896$	hochgradig signifikant
Hochschulkategorie	3,785e-14 (<0,001)	$\chi^2 (4) = 68,948$	hochgradig signifikant
Region	1,782e-05 (<0,001)	$\chi^2 (5) = 29,582$	hochgradig signifikant
Prüfungsstatus	0,0028 (<0,05)	$\chi^2 (1) = 8,954$	signifikant

Die Analyse der Fehlermagnitude deutet auf Folgendes hin: Wenn ein orthographischer Fehler im Lernertext auftritt, weisen erneut alle untersuchten Faktoren einen signifikanten Einfluss auf die Höhe der Fehlerwerte auf. Der Faktor *Diktattext* erweist sich hierbei

als hochgradig signifikant ($p < 0,001$). Dies bedeutet, dass der Faktor *Diktattext* nicht nur die Wahrscheinlichkeit beeinflusst, ob ein orthographischer Fehler auftritt, sondern darüber hinaus auch die Anzahl der Fehler mitbestimmt, sofern mindestens ein Fehler im Lernertext vorkommt. Der Diktattext wirkt sich also auf beiden Ebenen des Hurdle-Modells signifikant aus: sowohl auf die Fehlerwahrscheinlichkeit als auch auf die Fehlerhäufigkeit.

Auch die Faktoren *Hochschulkategorie* und *Region* bleiben in dieser Stufe hochgradig signifikant ($p < 0,001$). Der Faktor *Prüfungsstatus* bestätigt ebenfalls seine Relevanz als signifikanter Prädiktor ($p = 0,0028$).

Zusammenfassend lässt sich mittels des Hurdle-Modells feststellen, dass sich alle vier Faktoren – *Diktattext*, *Hochschulkategorie*, *Region* und *Prüfungsstatus* – in beiden Modellstufen als signifikante Prädiktoren erweisen. Sie beeinflussen sowohl die Wahrscheinlichkeit für das Auftreten eines orthographischen Fehlers als auch dessen Fehlerwert. Eine detailliertere Interpretation dieser Befunde folgt im nächsten Abschnitt.

5.1.2 Der Faktor *Diktattext*

Die inferenzstatistische Analyse mittels des Hurdle-Modells verdeutlicht, dass der Faktor *Diktattext* auf beiden Stufen der Fehlerproduktion eine maßgebliche Wirkung entfaltet.

In der ersten Modellstufe (Fehlerwahrscheinlichkeit) erweist sich der Diktattext als statistisch signifikant ($p = 0,0036$). Dies belegt, dass die spezifische linguistische Gestaltung der Prüfungstexte bereits die initiale Hürde zur Fehlerfreiheit beeinflusst. Die Wahrscheinlichkeit, einen Text vollkommen fehlerfrei zu produzieren, ist somit nicht zufällig verteilt, sondern hängt signifikant vom jeweiligen Text ab.

In der zweiten Modellstufe hingegen erweist sich der *Diktattext* als hochgradig signifikant ($p < 0,001$). Sobald die Hürde der Fehlerfreiheit überschritten ist, hat die spezifische sprachliche Beschaffenheit des Textes einen massiven Einfluss darauf, wie viele orthographische Fehler insgesamt gemacht werden. Dieser statistische Befund lässt sich qualitativ dadurch erklären, dass die Diktattexte von PGG-Prüfung 2017 und 2019 völlig unterschiedliche linguistische Schwerpunkte setzen. Je nachdem, welche orthographischen Phänomene in einem Text besonders häufig vorkommen, summieren sich die Fehlerwerte in unterschiedlichem Ausmaß. Tabelle 45 stellt die linguistischen Merkmale beider Texte systematisch gegenüber.

Tabelle 45: Vergleich der Diktattexte in den PGG-Prüfungen 2017 und 2019

	2019	2017
Diktattext	98 % der Bundesbürger sehen regelmäßig fern. Sie wollen sich am Abend vor dem Fernseher unterhalten und informieren. Der Alltag ist stressig, die Leute freuen sich auf das Wochenende und wollen sich ausruhen, ausschlafen oder sich mit der Familie treffen. Viele wünschen sich mehr Zeit für Hobbys, Sport und Freunde. Auch die Arbeit im Garten ist beliebt und hilft gegen Stress. Einer-	Durch verschiedene Arten von Belohnungen werden die Kinder zum Weiterlernen motiviert. Computer sind auch geduldiger als Lehrer, wenn man etwas nicht gleich kann. Außerdem lassen sich verschiedene Lehrinhalte durch die technischen Möglichkeiten ganz deutlich veranschaulichen. Aus diesen Gründen sollte man den Computer nicht von Anfang an ablehnen. Viele Eltern haben

(Fortsetzung Tabelle 45)

	2019	2017
	seits gibt es mehr Freizeitangebote, andererseits müssen die Menschen aber sparen. Immer mehr Deutsche gehen ins Schwimmbad, fahren Fahrrad und kochen zusammen. Freizeitvergnügen muss nicht immer Geld kosten.	zwar die Sorge, dass Kinder vor dem Computer vereinsamen, aber das kann man verhindern, indem man die Spiel- und Lerndauer vorher genau festlegt.
Wörterzahl	88	73
Nominalisierungen	Meist Konkreta oder etablierte Nomen: Arbeit, Fernseher	zum Weiterlernen, Belohnungen, Möglichkeiten, Lehrer
Silbengelenkschreibung	wollen, stressig, treffen, Hobbys, müssen, immer, kochen, zusammen	lassen
morph. Silbengelenkschreibung	Alltag, Stress, Schwimmbad, muss	sollte, kann
Silbeninitiales <h>	sehen, ausruhen, gehen, Fernseher	
Dehnungs-<h>	mehr, fahren, Fahrrad	Belohnungen, Lehrer, ablehnen, Lehrinhalte
Komposita	Bundesbürger, Fernseher, Alltag, Wochenende, Freizeitangebote, Schwimmbad, Fahrrad, Freizeitvergnügen	Lehrinhalte, Spiel- und Lerndauer
Fremdwörter	Hobbys, Sport, Stress	Computer, technisch
Durchschnittliche Satzlänge	11 Wörter pro Satz	14,6 Wörter pro Satz
Satzstruktur	Hauptsätze: 10 Nebensätze: 0	Hauptsätze: 6 Nebensätze: 3
Konjunktion	nebenordnend: 7 unterordnend: 0	nebenordnend: 2 unterordnend: 4

Besonders augenfällig sind die Abweichungen im Bereich der syntaktischen Komplexität. Der Diktattext von 2017 weist mit einer durchschnittlichen Satzlänge von 14,6 Wörtern eine deutlich höhere Komplexität auf als der Text von 2019 (11 Wörter pro Satz). Dieser Befund wird durch die Analyse der Satzstruktur untermauert: Während der Text von 2019 ausschließlich aus 10 einfachen Hauptsätzen besteht, enthält der Text von 2017 drei Nebensätze, die durch vier unterordnende Konjunktionen (*als, wenn, dass, indem*) eingeleitet werden. Diese verschachtelte Struktur stellte an die Prüflinge weitaus höhere Anforderungen hinsichtlich der Verarbeitung von Satzgefügen.

Qualitativ unterscheiden sich die Texte zudem signifikant in ihren orthographischen Schwerpunkten: Der Text von 2017 fordert die Lernenden besonders im Bereich der Nominalisierungen (z. B. die substantivierte Infinitivkonstruktion *zum Weiterlernen*) sowie bei der Dehnungs-<h>-Schreibung in Wörtern wie z. B. *Lehrer, ablehnen, Lehrinhalte*.

Im Gegensatz dazu liegt die primäre Herausforderung des Diktattexts von 2019 auf der lexikalischen und phonographischen Ebene. Der Text weist eine signifikant hohe Dichte an Silbengelenksschreibungen auf. Wörter wie *wollen*, *müssen*, *Stress*, *Schwimmbad* oder *zusammen* fordern die korrekte Markierung von Kurzvokalen durch Konsonantenverdopplung. Zudem verlangt die hohe Anzahl an Komposita (z. B. *Bundesbürger*, *Freizeitangebote*, *Freizeitvergnügen*) von den Lernenden ein orthographisches Detailwissen bei den morphologischen Schreibungen.

Dieser qualitative Unterschied, dass der Text von 2017 eine hohe **syntaktische Schwierigkeit** darstellt, während der Text von 2019 eine moderate **lexikalische Herausforderung** bietet, wird durch eine nähere Betrachtung in Bezug auf einzelne Fehlerkategorien dargestellt. Damit die Ergebnisse anschaulicher angezeigt werden, lassen sich in der Abbildung 33 Fehlerkategorien zu fünf Oberkategorien zusammenführen: phonographische Fehler (PHO: inkl. [Graph_PGK], [Graph_BF], [Graph_VQM], [Graph_S-Schreibung], [Graph_SG]; SIL: [Silben_AR], [Silben_SK], [Silben_ER]; MOR: [Morph_Konstanz], [Morph_Affix], [Morph_Kompo]; SYN: [Syn_GKS], [Syn_GZS], [Syn_dass-das], [Syn_man-Mann] und SON: [Son_FW], [Son_Gra-Ab], [SON]).

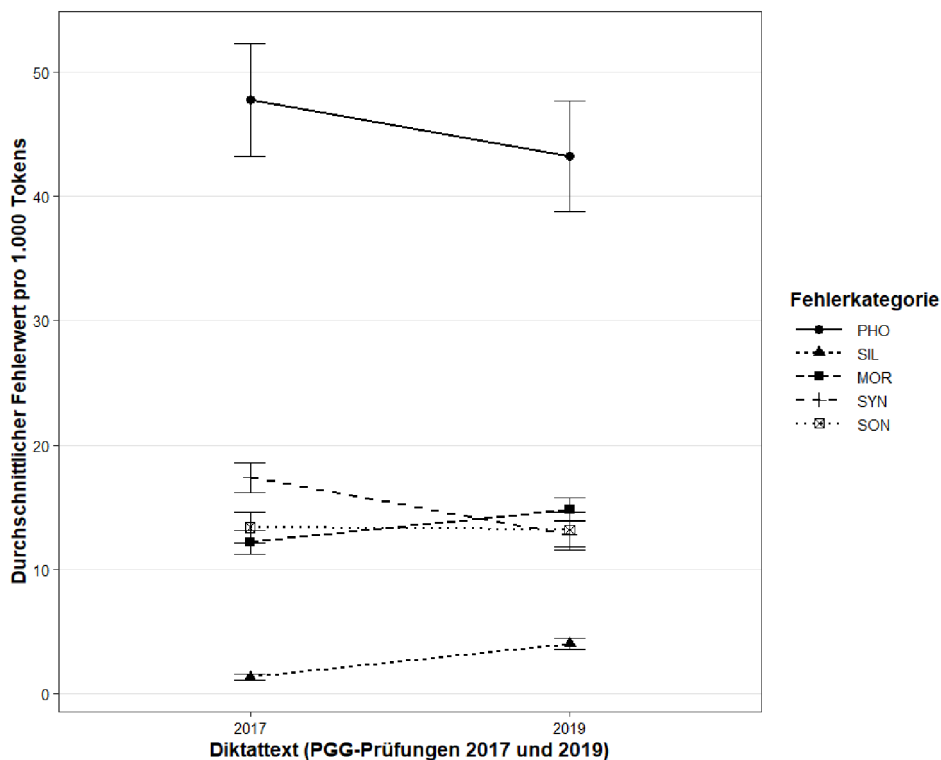


Abbildung 33: Durchschnittlicher Fehlerwert (inkl. Nullwerte) nach Diktattext (PGG-Prüfung 2017 und 2019) und Fehlerkategorie; Fehlerwert pro 1.000 Tokens; Fehlerbalken zeigen 1 SE

Auffällig ist der Unterschied in der Kategorie SYN, deren Fehlerwert im Jahr 2017 deutlich höher als derjenige im Jahr 2019 ist. Dies spiegelt die vorangegangene Textanalyse eindrücklich wider: Die hohe syntaktische Komplexität des Ausgangstexts mit seinen langen Sätzen und Nebensatzkonstruktionen führte bei den Lernenden erwartungsgemäß zu mehr Fehlern im syntaktischen Bereich (z. B. bei der Unterscheidung von *dass* und *das*). In den Kategorien SIL und MOR zeigt sich hingegen ein umgekehrtes Bild. Hier stieg der durchschnittliche Fehlerwert von 2017 nach 2019 an. Dies korrespondiert direkt mit der im Text 2019 beobachteten hohen Dichte an Silbengelenksschreibungen und Komposita, was die spezifischen Herausforderungen dieses Jahrgangs in den silbischen und morphologischen Bereichen bestätigt.

Der Diktattext beeinflusst einerseits signifikant die Wahrscheinlichkeit für das Auftreten eines orthographischen Fehlers und determiniert andererseits durch seine spezifische linguistische Struktur, in welchen Bereichen und in welchem Ausmaß diese orthographischen Fehler produziert werden. Die hochsignifikante Varianz in der zweiten Modellstufe ($p < 0,001$) ist somit das Resultat dieser unterschiedlichen linguistischen Profile: Die syntaktischen Schwierigkeiten von 2017 und die morphologisch-phonographischen Hürden von 2019 führen zu unterschiedlichen Gesamtschweregraden der fehlerhaften Texte.

5.1.3 Der Faktor *Region*

Die Ergebnisse des Hurdle-Modells verdeutlichen, dass der Faktor *Region* einen konsistenten und hochsignifikanten Einfluss auf die orthographische Fehlererscheinung hat. Dies legt nahe, dass der geografische Standort der Hochschule nicht nur darüber mitentscheidet, ob Lernende überhaupt einen orthographischen Fehler machen, sondern auch maßgeblich die Höhe der Fehlerwerte bestimmt, sobald solche Abweichungen auftreten.

Abbildung 34 vergleicht die durchschnittlichen Fehlerwerte von Studierenden aus unterschiedlichen Regionen. Die Fehlerwerte der Teilnehmenden aus Nordwest- und Südwestchina sind signifikant höher als die der Gruppen aus Zentral- und Südchina, Ostchina und Nordchina.

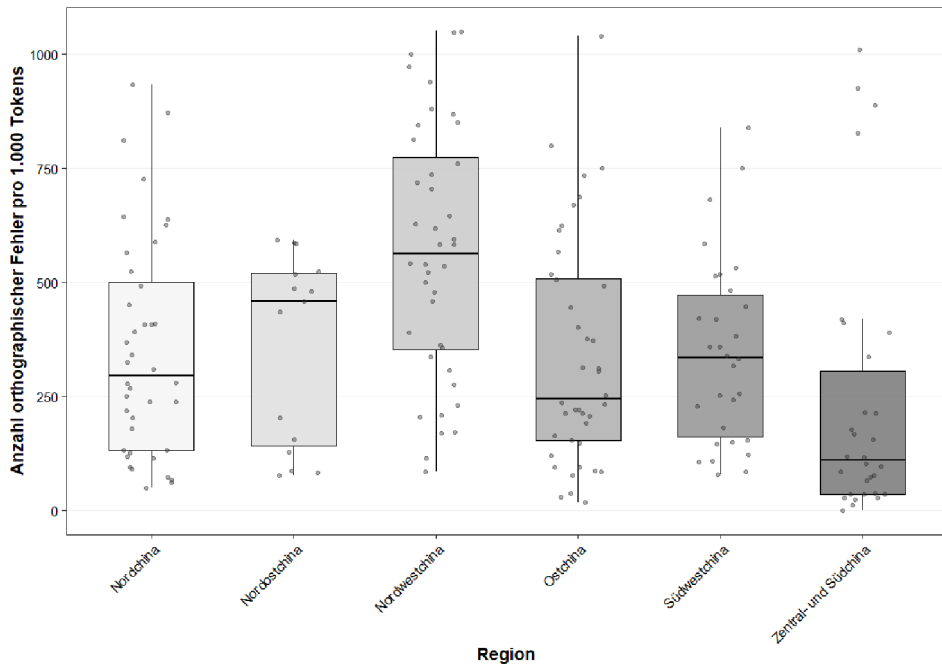


Abbildung 34: Variabilität der individuellen orthographischen Fehlerwerte pro Diktattext nach Regionen; jeder Punkt repräsentiert die Summe aller Fehlerstellen pro 1.000 Tokens eines einzelnen Diktattextes

Diese regionalen Unterschiede in der orthographischen Performanz korrelieren stark mit den sozioökonomischen und bildungspolitischen Gegebenheiten Chinas. Renommiertere Universitäten sind, wie das Shanghai Ranking, 2024 zeigt, überwiegend in den wirtschaftlich starken Regionen Nord- und Ostchinas (z. B. in Beijing, Shanghai, Zhejiang, Jiangsu, Anhui) verortet. Ein ähnliches Bild zeigt sich in der chinesischen Germanistik: Die führenden Studiengänge sind ebenfalls in Shanghai oder Beijing bzw. in den wirtschaftsstarken Küstenprovinzen angesiedelt (vgl. Marioulas & Wu, 2015, S. 39–41).

Neben dem Hochschulranking spielt der Standort eine entscheidende Rolle für die Studienwahl. Viele Studierende möchten ein Studium in einer wirtschaftlich hochentwickelten Metropole absolvieren, um nach dem Abschluss bessere Chancen auf dem Arbeitsmarkt zu haben. Dies gilt insbesondere für den Fachbereich für Fremdsprachen wie z. B. Germanistik, da sich die meisten Berufsfelder in Metropolen wie Beijing und Shanghai oder in anderen florierenden Städten wie Qingdao und Nanjing konzentrieren. Dies kann dazu führen, dass eine nominell niedriger eingestufte Hochschule (z. B. Kategorie B) in einer attraktiven Region auch Studieninteressierte mit sehr guten Gao-kao-Ergebnissen anzieht, die andernorts für eine Hochschule der Kategorie A zugelassen worden wären.

Diese regionalen Unterschiede beeinflussen nicht nur die Studierendenschaft, sondern auch die Lehrkräfte. Eine Anstellung an einer Hochschule in einer entwickelten Region bietet Germanist*innen in der Regel bessere finanzielle Konditionen und um-

fassendere Forschungsressourcen. Seit der Expansion der chinesischen Germanistik ab dem Jahr 2000 ist der Bedarf an qualifiziertem Lehrpersonal stark gestiegen. Da es in China bislang keine Institution gibt, die dezidiert Deutschlehrkräfte ausbildet, werden diese oft direkt aus dem Kreis promovierter Germanist*innen rekrutiert, die häufig keine spezifische didaktische Ausbildung erhalten haben (vgl. Zhao, 2020, S. 58). Für jene Lehrkräfte, die eine fachspezifische DaF-Lehrerbildung im Rahmen eines Austauschprogramms oder eines Studiums im Bereich Deutsch als Fremdsprache in Deutschland absolviert haben, sind die Hochschulen in den ökonomisch besser gestellten Regionen Chinas daher die attraktiveren Arbeitgeber.

Die Expansion der chinesischen Germanistik hat zwar auch zur Gründung zusätzlicher Institute in weniger entwickelten Provinzen geführt, doch weisen diese oft ungünstigere Studienbedingungen auf. In strukturschwächeren Regionen können die 211-Universitäten noch von einer Mitfinanzierung durch zentrale Ministerien profitieren und somit bestimmte Nachteile ausgleichen. Die Ausstattung anderer Hochschulen ist dagegen oft deutlich schlechter (vgl. Marioulas & Wu, 2015, S. 40). Beispielsweise haben Germanistikstudierende an Hochschulen in den westlichen Gebieten Chinas häufig einen wesentlich eingeschränkten Zugang zu Bildungsressourcen außerhalb des Unterrichts. Ein exemplarisches Beispiel für diese Diskrepanz sind die Bibliotheken der Germanistik-Fakultäten: Deren Bestände und Ausstattung können oft nicht mit denen in wirtschaftlich entwickelten Regionen mithalten.

Zusammenfassend lässt sich der signifikante Einfluss des Faktors *Region* auf die Fehlerwerte in Lernertexten als komplexer externer Faktor interpretieren. Dieser Faktor kann als ein Bündel von Kontrollparametern verstanden werden, welche den Zustand des dynamischen Lerner-Systems beeinflussen. Die Ergebnisse deuten darauf hin, dass Lernende in ressourcenstärkeren Regionen insgesamt eine höhere orthographische Kompetenz aufweisen, was sich in niedrigeren Fehlerwerten manifestiert.

5.1.4 Der Faktor Prüfungsstatus

In den beiden Stufen des Hurdle-Modells zeigt der Faktor *Prüfungsstatus* eine hochgradig signifikante Wirkung sowohl in Bezug auf die Wahrscheinlichkeit des Fehlerauftretens ($p < 0,001$) als auch auf die Höhe der Fehlerwerte in fehlerhaften Lernertexten ($p < 0,001$). Abbildung 35 veranschaulicht zunächst die durchschnittlichen Fehlerwerte nach Prüfungsstatus.

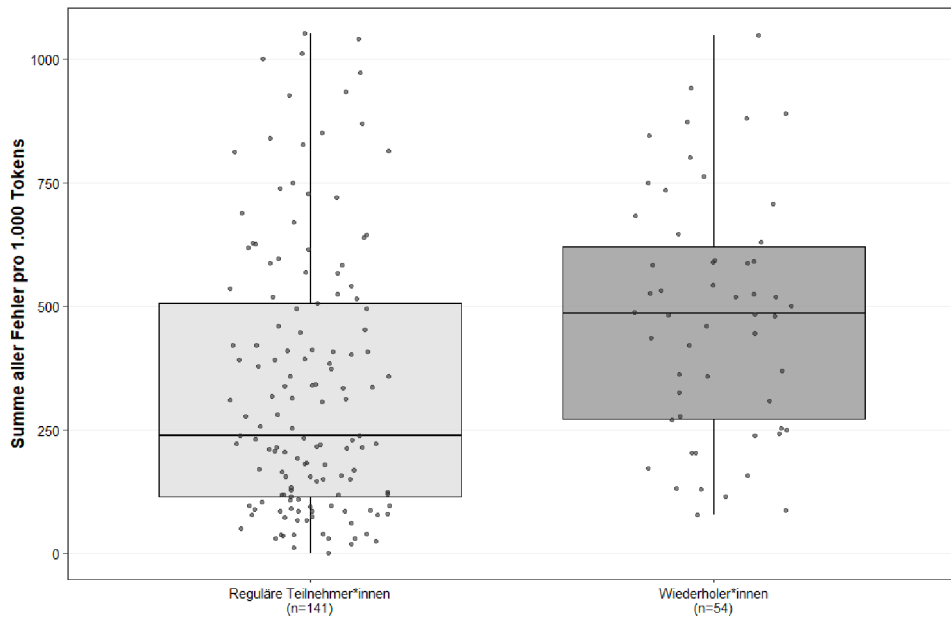


Abbildung 35: Variabilität der individuellen Fehlerwerte (inkl. Nullwerte) pro Diktattext nach Prüfungsstatus; jeder Punkt repräsentiert die Summe aller Fehlerstellen pro 1.000 Tokens eines einzelnen Diktattextes

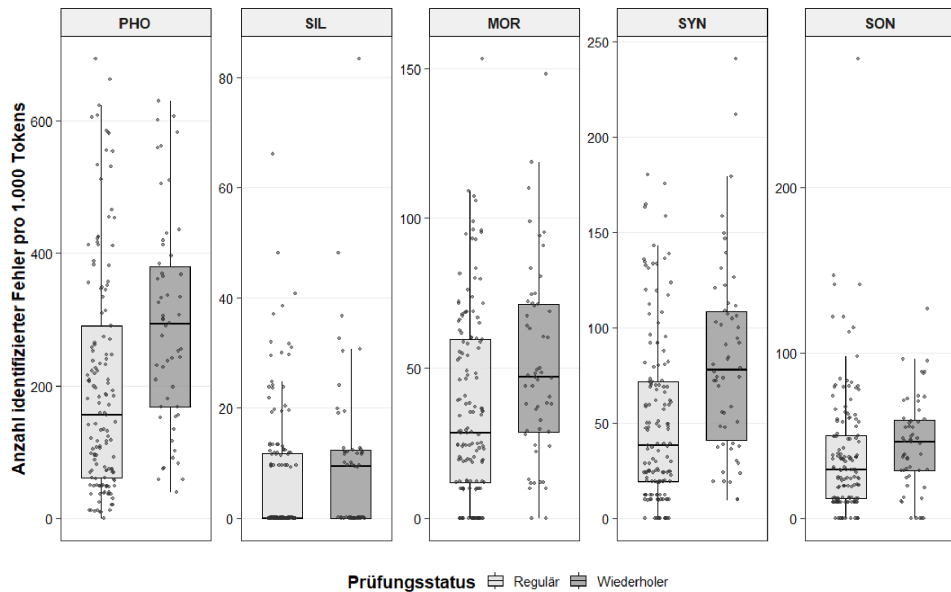


Abbildung 36: Differenzierte Fehleranalyse (inkl. Nullwerte) nach Oberkategorien und Prüfungsstatus; jeder Punkt repräsentiert die Summe der Fehlerstellen pro 1.000 Tokens in der jeweiligen Fehlerkategorie eines einzelnen Diktattextes

Obwohl die Wiederholer*innen in den dritten und späteren Studienjahren bereits zum zweiten Mal an der PGG-Prüfung teilnahmen und somit eine längere Vorbereitungszeit hatten, weisen sie einen signifikant höheren orthographischen Fehlerwert auf als die Lernenden im regulären Jahrgang. Dieses Ergebnis deutet darauf hin, dass dieser Leistungsunterschied nicht auf einen allgemeinen Rückgang des Lernfortschritts zurückzuführen ist, sondern auf die spezifische Zusammensetzung der Kohorte (Selektionseffekt). Die persistente Fehleranfälligkeit trotz längerer Lernzeit deutet auf verfestigte Schwierigkeiten im orthographischen Bereich hin, was eine genauere Betrachtung der Fehlerwerte nach Fehlerkategorien erfordert (vgl. Abbildung 36).

Eine detaillierte Aufschlüsselung der Fehlerwerte nach linguistischen Kategorien (vgl. Abbildung 36) bestätigt das Muster, das bereits in der Gesamtanalyse sichtbar wurde. Die Gruppe der Wiederholer*innen weist in nahezu allen Teilbereichen systematisch höhere Fehlerwerte auf als die reguläre Gruppe. Besonders deutlich wird dies in den Kategorien PHO und SYN, wo die Mediane der Wiederholer*innen (dunkelgraue Boxen) das Niveau der regulären Teilnehmenden (hellgraue Boxen) weit übersteigen. Auch in den Kategorien MOR und SON ist dieser Leistungsabstand signifikant. Lediglich im Bereich SIL fallen die Unterschiede geringer aus, was auf eine relativ stabile Beherrschung der silbischen Strukturen in beiden Gruppen hindeutet.

Grundsätzlich wird in der Spracherwerbsforschung diskutiert, inwiefern die Auseinandersetzung mit einem Schriftsystem zur Generierung impliziter Hypothesen führt. Gombert und Fayol (1992) konnten beispielsweise zeigen, dass Kinder bereits in der vorschulischen, voralphabetischen Entwicklungsphase die Prinzipien der Buchstabenstruktur eigenaktiv in Selbstlernprozessen entwickeln. Auch Bredel (2006) konnte für die Groß- und Kleinschreibung zeigen, dass Erstklässler bereits zu über 69 % korrekte Schreibungen vornehmen, auch ohne explizite Instruktion seitens der Lehrkraft.

Für den gesteuerten Fremdspracherwerb im chinesischen Hochschulkontext sind diese Erkenntnisse aus dem L1-Erwerb jedoch nur bedingt übertragbar. Während Kinder im Erstspracherwerb von einem reichen, multimodalen Input umgeben sind, ist der zielsprachliche Kontakt für die chinesischen Deutschlernenden weitgehend auf den universitären Unterricht beschränkt. Hierbei spielen die curricularen Rahmenbedingungen eine entscheidende Rolle.

Ein Blick auf den Studienverlaufsplan der Germanistikausbildung an der Anhui-Universität (übersetzt und zusammengestellt nach dem Studienverlaufsplan für das Jahr 2024; siehe Anhang 8.4) verdeutlicht eine signifikante Lücke: Das viersemestrige Grundstudium besteht fast ausschließlich aus grundlegendem Sprachunterricht und allgemeinbildenden Fächern. Im Hauptstudium (ab dem dritten Jahr) verlagert sich der Fokus auf fachspezifische und anwendungsorientierte Inhalte wie Übersetzung, Dolmetschen, Literatur- und Sprachwissenschaft. Das Thema Orthographie bleibt jedoch während des gesamten Studiums weitgehend ein Randthema. Eine systematische Auseinandersetzung mit den deutschen orthographischen Prinzipien findet weder in der grundlegenden Sprachausbildung noch im späteren Fachstudium statt. Diese curriculare Leerstelle im Bereich der Orthographie hat zur Folge, dass der Erwerb der deutschen Rechtschreibung weitgehend ohne systematische Anleitung erfolgen muss. Den Ler-

nenden bleibt infolgedessen kaum eine andere Wahl, als durch einen Prozess der Selbstorganisation eigene Hypothesen über die Regeln und Prinzipien des Schriftsystems zu entwickeln.

Bei der Interpretation der Fehlerentwicklung muss zudem berücksichtigt werden, dass es sich bei den vorliegenden Daten um eine Querschnittsanalyse handelt. Der beobachtete Anstieg der Fehlerwerte ab dem dritten Studienjahr ist nicht als Rückgang der Kompetenz zu missverstehen. Vielmehr könnte er auf einen Selektionseffekt zurückzuführen sein, da die Prüfungswiederholer*innen bereits im regulären Durchgang Schwierigkeiten aufwiesen. Um den tatsächlichen, individuellen Erwerbsverlauf jenseits solcher Selektionseffekte nachzuzeichnen, wird im folgenden Kapitel eine Analyse der Längsschnittdaten im Erwerbskorpus vorgenommen.

5.1.5 Der Faktor Hochschulkategorie

Bereits für den schulischen Bereich konnte in Studien wie der von Fix (2002) ein signifikanter Unterschied in der orthographischen Leistung zwischen Schüler*innen der Realschule und der Hauptschule nachgewiesen werden. Auch H.-G. Müller und Schroeder (2022) belegen einen signifikanten Einfluss der Schulart auf die Kompetenz in der Laut-Buchstaben-Zuordnung. Diese bildungsgangbezogenen Unterschiede in der orthographischen Kompetenz scheinen sich auf Hochschulebene fortzusetzen. So zeigen die Daten im DeChiLKo-Prüfungskorpus, dass der Faktor *Hochschulkategorie* erwartungsgemäß statistisch signifikant ist, sowohl hinsichtlich der orthographischen Fehlerwahrscheinlichkeit ($p < 0,001$) als auch der Fehlermagnitude (der Fehlerwerte in fehlerhaften Texten, $p < 0,001$).

Wie aus der Abbildung 37 ersichtlich ist, weisen Studierende von Hochschulen der Kategorien A und B im Durchschnitt niedrigere Fehlerwerte auf als jene aus den Kategorien D und E. Die höchsten orthographischen Fehlerwerte zeigen jedoch die Studierenden von Hochschulen der Kategorie C, was mit insgesamt schwächeren Prüfungsergebnissen in dieser Gruppe korrespondiert (vgl. Abbildung 38). Die Verteilung der Gesamtnoten nach Hochschulkategorie macht deutlich, dass Studierende der Kategorie-A-Hochschulen tendenziell höhere Gesamtnoten in der PGG-Prüfung erzielen. Die niedrigsten Noten finden sich wiederum in Kategorie C, gefolgt von D und E.

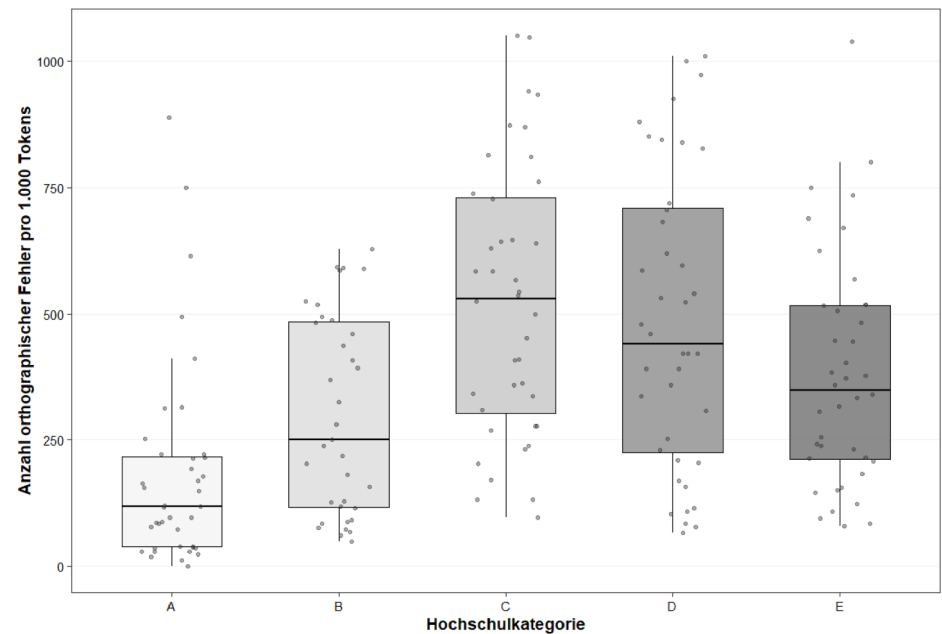


Abbildung 37: Variabilität der individuellen Fehlerwerte (inkl. Nullwerte) pro Diktattext nach Hochschulkategorie; jeder Punkt repräsentiert die Summe aller Fehlerstellen pro 1.000 Tokens eines einzelnen Diktattextes

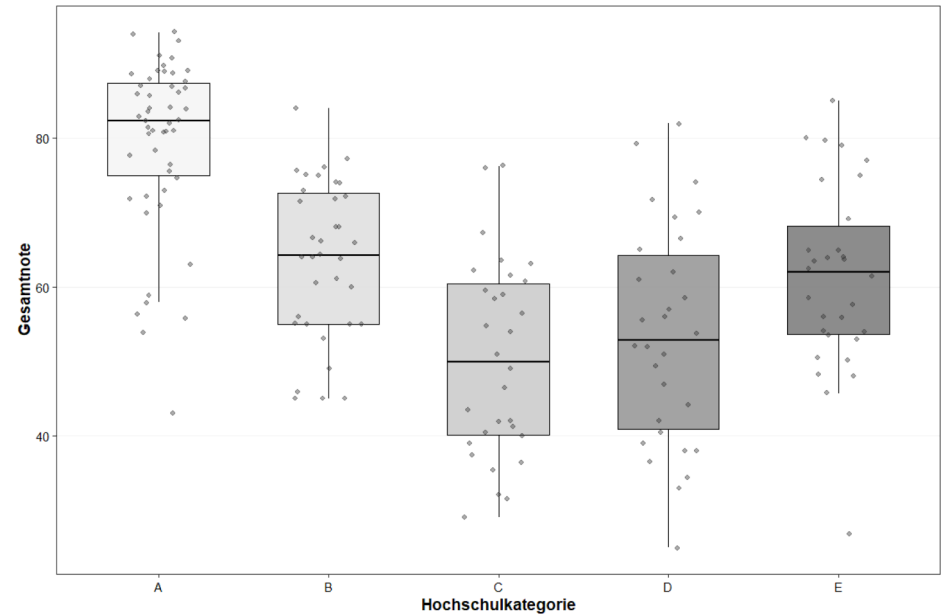


Abbildung 38: Verteilung der PGG-Gesamtnoten nach Hochschulkategorie (A–E); Einzelwerte als Datenpunkte

Ein möglicher Erklärungsansatz für diese Befunde ergibt sich aus der leistungsselektiven Struktur des chinesischen Hochschulsystems. Der Zugang zu renommierten Hochschulen (z. B. Kategorien A und B) ist in der Regel an hohe *Gaokao*-Ergebnisse gebunden. Viele empirische Studien belegen, dass die *Gaokao* ein valider Prädiktor für akademische Leistungsfähigkeit ist und signifikant mit dem späteren Studienerfolg korreliert (vgl. Bai, Chi & Qian, 2014; Farley & Yang, 2020). Damit fungiert die Hochschulkategorie indirekt als ein Indikator für die kognitive Leistungsfähigkeit der Studierenden. Dass kognitive Leistungsfähigkeiten eine zentrale Rolle für den Erwerb schriftsprachlicher Kompetenzen spielen, wird durch mehrere empirische Untersuchungen gestützt: Corvacho Del Toro (2013, S. 194) konnte eine starke Korrelation zwischen der Grundintelligenz von Kindern und ihrer Rechtschreibleistung nachweisen. Auch Becker (2013, S. 228) identifizierte insbesondere bei einsprachig aufwachsenden Kindern einen deutlichen Zusammenhang zwischen kognitiver Leistungsfähigkeit und orthographischer Kompetenz. In ähnlicher Weise hebt Bulut (2018, S. 110) hervor, dass kognitive Faktoren nicht nur den stärksten Effekt auf die korrekte Schreibung des Schwa-Lauts ausüben, sondern auch die Schreibung anderer Rechtschreibphänomene signifikant beeinflussen. Hochschulen der unteren Kategorien (C, D, E) rekrutieren folglich häufig Studierende, die im Durchschnitt über geringere sprachliche und kognitive Eingangsvoraussetzungen verfügen. Dies könnte die beobachteten Unterschiede in der orthographischen Kompetenz mitbedingen.

Die Verteilung der orthographischen Fehlerwerte nach Hochschulkategorien entspricht zwar im Großen und Ganzen der Verteilung der Gesamtnoten, jedoch sind die Unterschiede zwischen den Kategorien weniger ausgeprägt als bei den Gesamtnoten. Dies lässt sich durch mehrere Faktoren erklären: Zum einen besteht, wie in Abschnitt 5.1.3 diskutiert, in China eine erhebliche regionale Disparität im Hochschulwesen, sodass *Gaokao*-Punktzahlen und das formale Hochschulranking nicht immer in einer streng linearen Beziehung zueinanderstehen. Zum anderen sind die meisten Studierenden zu Beginn ihres Germanistikstudiums sprachlich Nullanfänger*innen (vgl. Zhao, 2020, S. 54). Dies bedeutet, dass der spätere Lernerfolg und die orthographische Kompetenz nicht nur von den Eingangsvoraussetzungen, sondern auch stark von der Qualität der universitären sprachlichen Ausbildung beeinflusst werden. Wie bereits dargestellt, korrelieren auch Lehrkräfte und infrastrukturelle Bedingungen systematisch mit der Hochschulkategorie, was sich wiederum auf den Lernverlauf und die Entwicklung des orthographischen Systems der Lernenden auswirken kann.

In dem folgenden Kapitel wird der Fokus auf die tatsächliche Entwicklung der orthographischen Kompetenz im Zeitverlauf gelegt. Mithilfe der Längsschnittdaten aus dem DeChiLKo-Erwerbskorpus wird untersucht, wie sich die Orthographieleistung von fünf chinesischen Germanistikstudent*innen während der ersten drei Semester ihres Grundstudiums entwickelt und welchen Einfluss Faktoren wie die Lerndauer und die Aufgabentypen auf die orthographischen Fehlerwerte ausüben.

5.2 Einflussfaktoren für die Orthographiekompetenz im Erwerbskorpus

Im DeChiLKo-Erwerbskorpus wurden Diktattexte von fünf Germanistikstudierenden longitudinal über ihre ersten drei Semester (Wintersemester 2021/22 bis zum Wintersemester 2022/23) an der Anhui-Universität erhoben. Die Anhui-Universität befindet sich in Ostchina und gehört zur Hochschulkategorie A. Für diese fünf Studierenden gilt Deutsch als ihre zweite Fremdsprache. Sie haben alle ab dem 3. Schuljahr in der Grundschule Englisch als erste Fremdsprache gelernt. Das Deutschlernen begann mit dem Eintritt in das Bachelorstudium im September 2021. Die Lerndauer wurde zu jedem Erhebungszeitpunkt in Monaten berechnet, wobei als Beginn des Deutschlernens der 1. September 2021 angesetzt wurde. Als eine klassische Übungsform wurde das Diktat bereits ab dem ersten Monat im Unterricht des Grundstudiums an der Germanistikabteilung der Anhui-Universität eingesetzt. Diese Übungsform wird bis zum Ende des zweiten Studienjahres beibehalten und dient auch als wichtiges Training für die PGG-Prüfung.

Die im Erwerbskorpus gesammelten Diktattexte wurden entweder als regelmäßige, unbenotete Unterrichtsübungen im Rahmen der Fächer „Grundstufe Deutsch“ oder „Hörverstehen Deutsch“ geschrieben oder unter Prüfungsbedingungen als Teil von Semesterprüfungen, wo sie ähnlich wie in der PGG-Prüfung als Teil des Hörverstehens bewertet wurden. Dieser Unterschied wurde als Faktor Aufgabentyp erfasst und in die Modellierung einbezogen. Der Faktor Diktattext selbst wurde ursprünglich ebenfalls als potenzieller Prädiktor in Betracht gezogen. Aufgrund einer sehr hohen Kollinearität mit dem Faktor Aufgabentyp, die zu Redundanz in der Modellmatrix führte, wurde der Faktor Diktattext aus dem finalen Modell entfernt, um die Stabilität und eindeutige Interpretierbarkeit der Modelleffekte zu gewährleisten.

5.2.1 Statistische Analyse mittels Hurdle-Modells

Zur Analyse des orthographischen Fehlerwertes im Erwerbskorpus wurde, wie schon beim Prüfungskorpus, das Hurdle-Modell (siehe Kapitel 5.1.1) gewählt. Damit können das Fehlerauftreten und die Fehlerausprägung (Magnitude) getrennt voneinander modelliert werden.

In beiden Modellstufen wurden die Faktoren Aufgabentyp, Lerndauer und die 18 Fehlerkategorien als feste Effekte (Kontrollfaktoren) berücksichtigt. Die fünf Teilnehmenden wurden hingegen als zufällige Effekte modelliert, um die Varianz zwischen den einzelnen Lernenden zu kontrollieren und somit ihre individuellen Lernpfade zu berücksichtigen. Wie bei der Modellierung der Faktoren im Prüfungskorpus wurden die p-Werte der einzelnen Faktoren im Erwerbskorpus ebenfalls für beide Modellstufen mittels Likelihood-Quotienten-Tests bestimmt. Dabei wurde jeweils das volle Modell mit einem reduzierten Modell verglichen, aus dem der zu testende Faktor entfernt wurde (vgl. Winter, 2020, S. 260–261). Die folgende Tabelle fasst die Ergebnisse der Likelihood-Quotienten-Tests zusammen.

Tabelle 46: Ergebnisse der Likelihood-Quotienten-Tests in den beiden Stufen des Hurdle-Modells im Erwerbskorpus

Hurdle-Modell	Faktor	p-Werte	Chi-Quadrat (Freiheitsgrade)	Signifikanz
Stufe I (Fehlerauftreten)	Fehlerkategorie	$< 2,2e-16$	$\chi^2 (16) = 797,99$	hochgradig signifikant
	Lerndauer	$< 2,2e-16$	$\chi^2 (1) = 94,055$	hochgradig signifikant
	Aufgabentyp	0,060	$\chi^2 (1) = 3,533$	nicht signifikant
Stufe II (Höhe der Fehlerwerte)	Fehlerkategorie	$< 2,2e-16$	$\chi^2 (16) = 160,77$	hochgradig signifikant
	Lerndauer	0,003 ($< 0,05$)	$\chi^2 (1) = 8,762$	signifikant
	Aufgabentyp	0,2377	$\chi^2 (1) = 1,4608$	nicht signifikant

Die Ergebnisse der statistischen Analyse bekunden, dass die Faktoren *Fehlerkategorie* und *Lerndauer* in den beiden Stufen des Hurdle-Modells eine statistische Signifikanz ($p < 0,05$) aufweisen. Dies bedeutet, dass die Lerndauer einen signifikanten Einfluss sowohl darauf hat, ob überhaupt orthographische Fehler einer bestimmten Kategorie auftreten, als auch darauf, wie hoch die Fehlerwerte in fehlerhaften Lernertexten ausfallen. Der hochsignifikante Einfluss der Lerndauer bestätigt die Erwartung, dass mit zunehmender Lernzeit eine Entwicklung der orthographischen Kompetenz stattfindet. Dieses Ergebnis wird im folgenden Abschnitt detaillierter im Hinblick auf die Entwicklungsverläufe einzelner Fehlerkategorien analysiert.

Der Faktor *Aufgabentyp* zeigt hingegen in beiden Modellstufen keinen signifikanten Einfluss. Dies deutet darauf hin, dass die orthographische Performanz der Lernenden nicht signifikant davon beeinflusst wird, ob ein Diktat als reguläre, unbenotete Unterrichtsübung oder als benotete Prüfungsaufgabe geschrieben wird. Auch wenn man annehmen könnte, dass sich die Lernenden während der Unterrichtsübung in einer weniger stressigen Situation befinden, zeigen sie keine systematische Verbesserung oder Verschlechterung ihrer orthographischen Leistung. Dies könnte darauf hindeuten, dass die orthographische Kompetenz ein relativ stabiles System darstellt, dessen Performanz weniger von situativen Faktoren wie Prüfungsdruck beeinflusst wird.

5.2.2 Der Faktor Lerndauer

Nachdem die Untersuchung des *Prüfungsstatus* bereits gezeigt hat, dass die orthographische Fehlererscheinung keinen rein stetigen Verlauf nimmt, rückt nun die Lerndauer als prozessualer Faktor in den Mittelpunkt. Anhand der longitudinalen Daten des Erwerbskorpus soll analysiert werden, wie sich die Fehlererscheinungen über die Zeit hinweg entwickeln. Hierfür werden die Fehlerkategorien zu fünf Oberkategorien zusammengefasst: [PHO] (phonographische Abweichungen), [SIL] (silbische Abweichungen), [MOR] (morphologische Abweichungen), [SYN] (syntaktische Abweichungen) sowie [SON] (sonstige Abweichungen).

Abbildung 39 zeigt die Variabilität und Entwicklung der orthographischen Fehlerwerte (pro 1.000 Tokens) für die fünf Oberkategorien über den Beobachtungszeitraum der ersten drei Semester. Durch die Einbeziehung der Nullwerte bildet diese Darstellung die globale Fehlerverteilung im Erwerbskorpus ab.

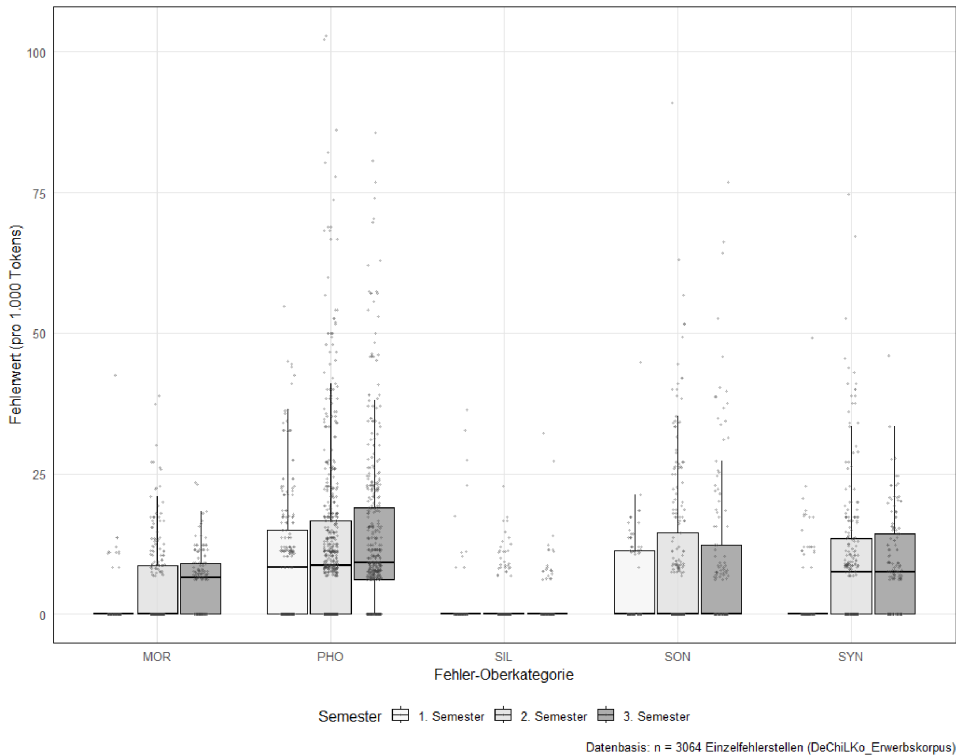


Abbildung 39: Veränderung der Fehlerwerte (inkl. Nullwerte) nach Oberkategorie und nach Semestern (Semester 1, 2, 3 von links); Einzelwerte als Datenpunkte.

Aus der grafischen Darstellung wird der dynamische Charakter des Orthographieerwerbs ersichtlich. Die Fehlerprofile der Lernenden sind nicht statisch, sondern weisen im Zeitverlauf charakteristische Verschiebungen auf. Zu Beginn des Studiums liegen die primären Schwierigkeiten im phonographischen Bereich (PHO), während die übrigen Kategorien durch eine hohe Dichte an Nullwerten geprägt sind, was auf eine noch begrenzte Anwendung komplexerer sprachlicher Strukturen hindeutet. Mit dem Übergang in das zweite Semester vollzieht sich eine markante Verschiebung in der Fehlerverteilung, die durch eine deutliche Zunahme der Variabilität und Frequenz insbesondere in der syntaktischen Kategorie (SYN) gekennzeichnet ist. Diese Entwicklung ist als Phase der systemischen Restrukturierung zu interpretieren, in der die Anwendung neuer sprachlicher Hypothesen zu einer temporären Destabilisierung des Gesamtsystems führt, während morphologische und silbische Fehlererscheinungen

tendenziell in den Hintergrund treten. Im dritten Semester mündet dieser Prozess in eine zunehmende Konsolidierung, wobei das Fehlerprofil im Vergleich zum Vorsemerster weitgehend stabil bleibt und die PHO weiterhin als die persistenteste Herausforderung mit der höchsten Individualität identifiziert werden kann.

Um die reine Fehlerintensität zu isolieren und den Einfluss der fehlerfreien Kategorien (Nullwerte) auszuschließen, kontrastiert Abbildung 40 die Datenbasis exklusive der Nullwerte in der jeweiligen Fehlerkategorie. Diese Darstellung entspricht der Stufe II des Hurdle-Modells und verdeutlicht die Höhe der Fehlerwerte, sofern diese auftreten.

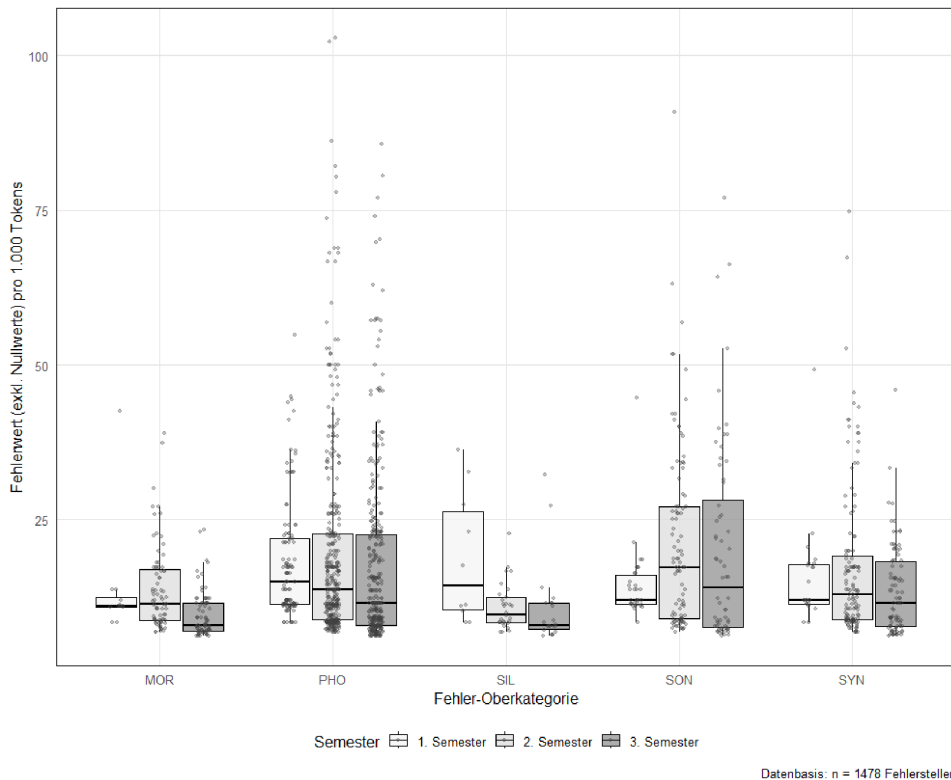


Abbildung 40: Abbildung 40 Veränderung der Fehlerwerte (exkl. Nullwerte) nach Oberkategorie und nach Semestern (Semester 1, 2, 3 von links); Einzelwerte als Datenpunkte.

Um die beobachteten Unterschiede der orthographischen Fehlererscheinungen zwischen den Semestern inferenzstatistisch zu überprüfen, wurden für jede Oberkategorie separate lineare gemischte Modelle berechnet. Die Struktur dieses Modells glich dem in Kapitel 5.2.1 beschriebenen Modell in der Stufe II des Hurdle-Modells, jedoch ohne den Faktor Fehlerkategorie.

Die Ergebnisse des Post-hoc-Tests mit dem R-Paket *emmeans* (Lenth, 2025) sind in der folgenden Tabelle 47 zusammengefasst. Der Schätzwert (Estimate-Wert) gibt die

Differenz der geschätzten mittleren Fehlerwerte zwischen den beiden verglichenen Semestern an. Ein positiver Wert bedeutet, dass im zuerst genannten Semester mehr Fehler gemacht wurden als im zweiten; ein negativer Wert bedeutet das Gegenteil.

Tabelle 47: Ergebnisse des Post-hoc-Tests zum Vergleich der durchschnittlichen Fehlerwerte (exkl. Nullwerte) zwischen Semestern nach Fehlerkategorie

Vergleich	Fehlerkategorie	Estimate-Werte	p-Werte
1. Semester – 2. Semester	PHO	0,149	0,122
2. Semester – 3. Semester		0,198	0,0096 (<0,05)
1. Semester – 3. Semester		0,347	0,0002 (<0,001)
1. Semester – 2. Semester	SIL	0,435	0,0121 (<0,05)
2. Semester – 3. Semester		0,125	0,6035
1. Semester – 3. Semester		0,560	0,0039 (<0,05)
1. Semester – 2. Semester	MOR	0,057	0,9159
2. Semester – 3. Semester		0,329	0,0002 (<0,001)
1. Semester – 3. Semester		0,386	0,0261 (<0,05)
1. Semester – 2. Semester	SYN	–0,002	0,999
2. Semester – 3. Semester		0,181	0,087
1. Semester – 3. Semester		0,179	0,388
1. Semester – 2. Semester	SON	–0,254	0,087
2. Semester – 3. Semester		0,095	0,646
1. Semester – 3. Semester		–0,159	0,473

Die statistische Analyse bestätigt die deskriptiven Beobachtungen. Ein Vergleich der Fehlerwerte zwischen dem ersten Semester und dem dritten Semester zeigt, dass die Lernenden innerhalb der ersten drei Semester signifikante Fortschritte in den phonographischen ($p < 0,001$), silbischen ($p < 0,05$) und morphologischen ($p < 0,05$) Bereichen machen. Die Entwicklung ist jedoch nicht linear. So findet beispielsweise eine signifikante Verbesserung im silbischen Bereich (SIL) bereits zwischen dem ersten und zweiten Semester statt ($p < 0,05$). In den phonographischen (PHO) und morphologischen (MOR) Bereichen hingegen zeigt sich ein signifikanter Fortschritt erst zwischen dem zweiten und dritten Semester ($p < 0,05$ bzw. $p < 0,001$). Dieses gestaffelte Entwicklungsmuster, bei dem verschiedene Subsysteme zu unterschiedlichen Zeiten an Stabilität gewinnen, ist ein typisches Merkmal dynamischer Systeme.

Im syntaktischen Bereich (SYN) ist über den gesamten Beobachtungszeitraum hinweg kein signifikanter Fortschritt zu verzeichnen. Eine differenzierte Analyse der Subkategorien in Abbildung 41 verdeutlicht, dass die Herausforderung primär in den

Bereichen Getrennt- und Zusammenschreibung (GZS) sowie Groß- und Kleinschreibung (GKS) vorkommt, während die dass/das-Schreibung eine untergeordnete Rolle mit nahezu konstanten Nullwerten einnimmt. Auffällig ist die Entwicklung bei der Getrennt- und Zusammenschreibung, die nach einem fast fehlerfreien ersten Semester im zweiten Semester einen massiven Anstieg der Medianwerte verzeichnet und auch im dritten Semester auf einem hohen Niveau bleibt. Bei der Groß- und Kleinschreibung zeigt sich ein charakteristischer nicht linearer Verlauf: Hier ist ein Anstieg der Fehlerwerte im zweiten Semester (Mittelwert ca. 10) zu beobachten, gefolgt von einer leichten Zunahme im dritten Semester.

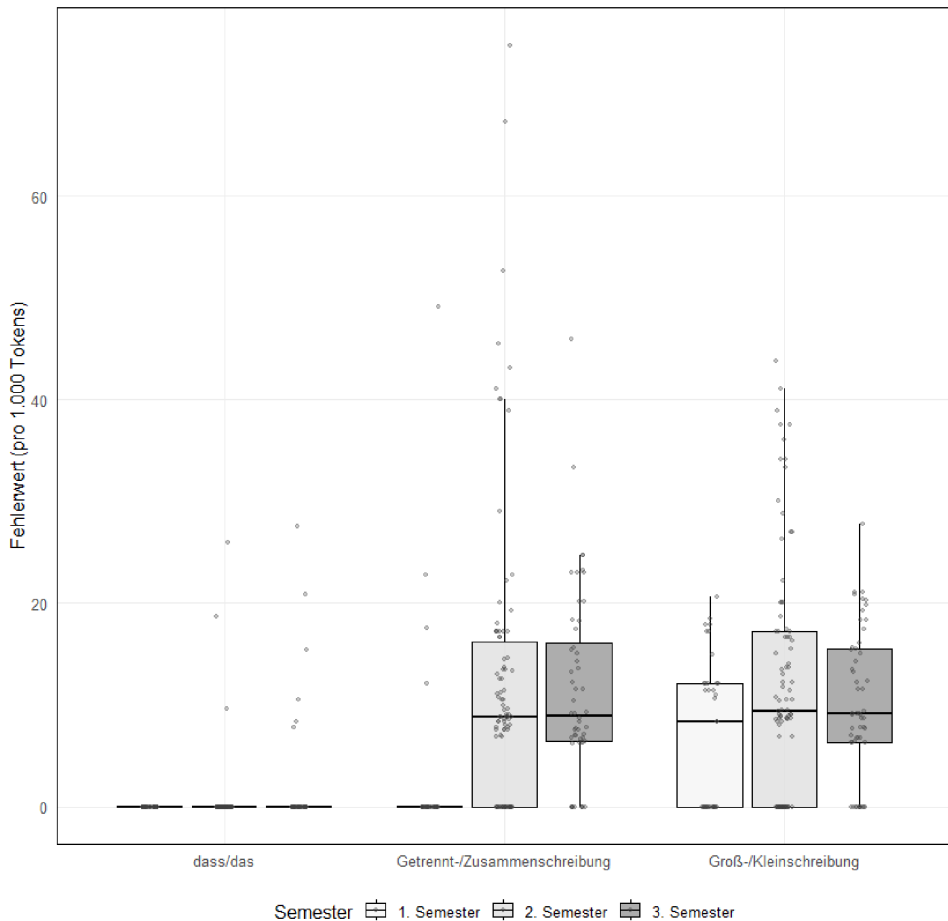


Abbildung 41: Variabilität der Fehlerwerte (inkl. Nullwerte) in den syntaktischen Kategorien über drei Semester; Einzelwerte als Datenpunkte.

Diese nicht lineare Entwicklung bei der Groß- und Kleinschreibung kann als Indiz für eine Restrukturierung des Wissenssystems interpretiert werden. Möglicherweise führt

die Auseinandersetzung mit komplexeren syntaktischen Strukturen und Nominalisierungsprozessen im zweiten Semester zu einer temporären Destabilisierung bereits bestehender, aber noch fragiler Hypothesen zur Großschreibung, was sich in einem erhöhten Fehlerwert und dem weiten Spektrum der Einzeldatenpunkte niederschlägt.

Für die Kategorie Sonstiger Bereich (SON) ist über den gesamten Beobachtungszeitraum hinweg zwar kein global signifikanter Unterschied zwischen den Semestern zu verzeichnen, jedoch stellt sie die einzige Oberkategorie dar, deren Fehlerwerte über die ersten drei Semester hinweg eine tendenzielle Zunahme aufweisen. Eine detaillierte Darstellung in Abbildung 42 deutet darauf hin, dass diese Entwicklung mit der Zunahme grammatisch erklärbarer Abweichungen, wie Deklinations- und Konjugationsfehlern, im zweiten und dritten Semester getragen wird. Während die Fehlerwerte im Bereich der Fremdwörter über alle Semester hinweg auf einem konstant niedrigen Niveau mit minimaler Streubreite verharren, zeigt die Subkategorie *Grammatisch* ab dem zweiten Semester eine signifikante Expansion sowohl der Medianwerte als auch des Spektrums der Einzeldatenpunkte.

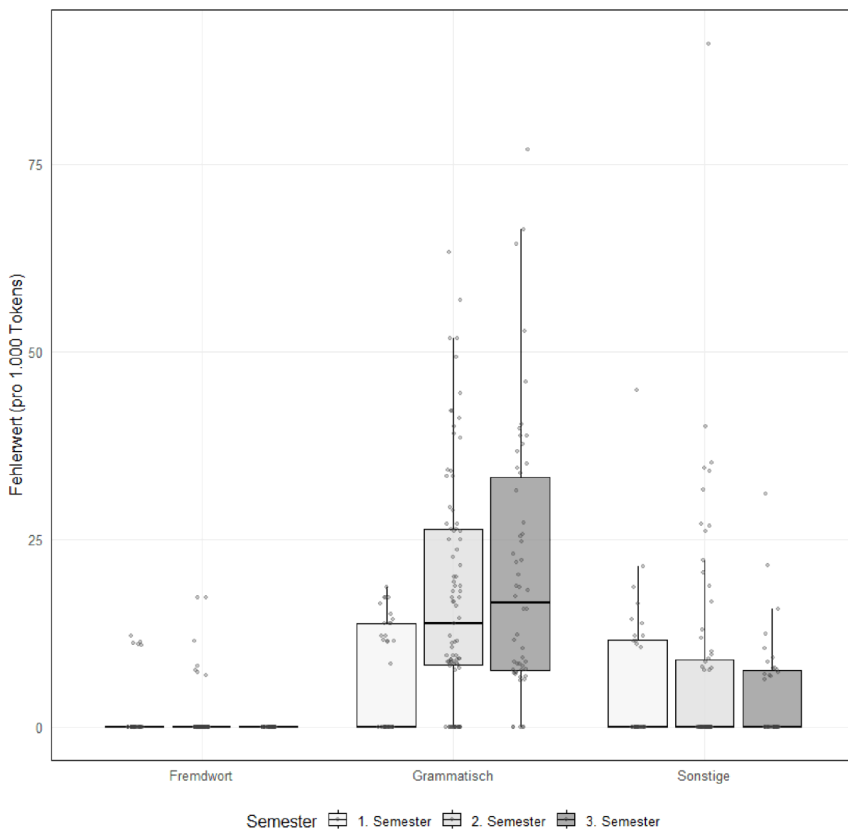


Abbildung 42: Variabilität der Fehlerwerte (inkl. Nullwerte) in den sonstigen Kategorien über drei Semester; Einzelwerte als Datenpunkte.

Dieses Phänomen ist besonders aussagekräftig vor dem Hintergrund der L1 Chinesisch, einer isolierenden Sprache ohne Flexion. Für die Lernenden stellt der Erwerb der deutschen Deklination und Konjugation eine erhebliche Herausforderung dar (vgl. Lay, 2017, S. 35).

Es darf jedoch nicht außer Acht gelassen werden, dass die Lernenden ihre individuellen Entwicklungspfade aufweisen, wie die facettierte Darstellung in Abbildung 43 signalisiert, wobei sich sowohl die Entwicklungsverläufe als auch die Fehlerverteilungen innerhalb der orthographischen Kategorien bei den fünf Probanden des Erwerbskorpus deutlich unterscheiden.

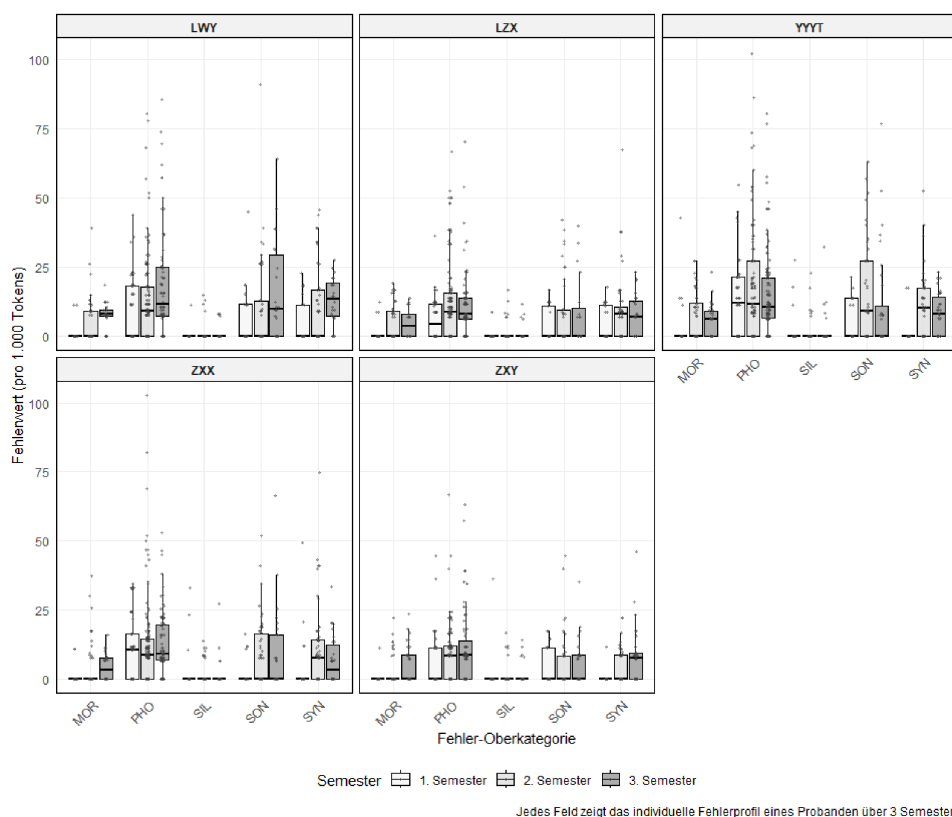


Abbildung 43: Individuelle Fehlerprofile der fünf Probanden nach Oberkategorien und Semestern; Einzelwerte als Datenpunkte

Beispielsweise zeigt es bei der Lernerin LWY ein Entwicklungspfad stetig steigender Fehlerwerte in fast allen Bereichen, wobei besonders in der Phonographie und im sonstigen Bereich zum dritten Semester hin eine erhebliche Ausweitung der individuellen Streubreite zu beobachten ist. Im Gegensatz dazu ist die Entwicklung bei LZX durch einen Zwischengipfel im zweiten Semester geprägt, bei dem vor allem die syntaktischen Fehler ihren Höchstwert erreichen und danach wieder leicht abflachen. Bei

YYT nehmen die Fehler im sonstigen und syntaktischen Bereich im zweiten Semester extrem stark zu, bevor sie im dritten Semester wieder deutlich sinken, was als typisches Indiz für eine systemische Restrukturierung zu interpretieren ist. Bei ZXX ist hingegen eine ausgeprägte „Umkehr-U-Kurve“ im syntaktischen Bereich zu verzeichnen, wobei die Boxplots verdeutlichen, dass die Phonographie über alle drei Semester hinweg eine sehr konstante, aber breit gestreute Fehlerquelle bleibt. Im Vergleich zum restlichen Kollektiv weist der Lernende ZXY die niedrigsten Fehlerwerte auf, zeigt jedoch nach einem sehr flachen Start zum Ende des Beobachtungszeitraums einen kräftigen Anstieg der Fehlerwerte, vor allem in den syntaktischen und morphologischen Kategorien. Diese Einzelfallanalysen unterstreichen, dass die orthographische Entwicklung nicht als linearer Fortschritt, sondern als Bewegung zwischen verschiedenen Attraktorzuständen zu verstehen ist, die durch individuelle Lernbedingungen moderiert werden.

Zusammenfassend zeigt die longitudinale Analyse des DeChiLKo-Erwerbskorpus, dass sich die orthographische Kompetenz der chinesischen Deutschlernenden in den ersten drei Semestern signifikant entwickelt, dieser Prozess jedoch hochgradig nicht linear und bereichsspezifisch verläuft. Während grundlegende phonographische und silbische Kompetenzen relativ früh verbessert werden, kommt es in komplexeren syntaktischen und grammatiknahen Bereichen zu Stagnationen oder sogar temporären Regressionen. Diese beobachtete Entwicklungsdynamik bestätigt zentrale Annahmen der CDST. Demnach ist der Orthographieerwerb kein rein kumulativer, sondern ein von kontinuierlicher Restrukturierung und individueller Variabilität geprägter Prozess.

6 Ergebnisse und Diskussion

6.1 Zusammenfassung und Diskussion der Ergebnisse

Die vorliegende Arbeit verfolgte das Ziel, die orthographische Kompetenz chinesischer Deutschlernender empirisch fundiert zu untersuchen. Im Mittelpunkt standen dabei folgende Forschungsfragen:

- Forschungsfrage 1: Welche spezifischen orthographischen Herausforderungen und typischen Fehlermuster ergaben sich bei chinesischen DaF-Lernenden beim Erwerb der deutschen Rechtschreibung?
- Forschungsfrage 2: Wie entwickelt sich die Rechtschreibkompetenz chinesischer Deutschlernender in der Anfangsphase des Deutschlerwerbs?
- Forschungsfrage 3: Welche Faktoren beeinflussen die beobachtbare orthographische Leistung chinesischer DaF-Lernender, und wie ist deren dynamisches Zusammenspiel zu charakterisieren?

Zur Beantwortung dieser Fragen wurde das Deutsch-Chinesische Lernerkorpus (DeChiLKo) erstellt, das 329 Diktattexte (31.674 Tokens) von chinesischen Germanistikstudierenden aus einem Prüfungskorpus (Diktat in den PGG-Prüfungen 2017 und 2019) und einem longitudinalen Erwerbskorpus (Diktat von fünf Proband*innen in den ersten drei Semestern) umfasst. Die orthographischen Abweichungen wurden mithilfe eines detaillierten, mehrschichtigen Annotationsschemas erfasst, das phonographische, silbische, morphologische, syntaktische und sonstige Aspekte berücksichtigt.

6.1.1 Spezifische orthographische Herausforderungen und Transferphänomene

Die erste Frage bezog sich darauf, anhand der vorliegenden lernersprachlichen Daten einerseits spezifische orthographische Herausforderungen und typische Fehlermuster chinesischer Deutschlernender in den unterschiedlichen Bereichen sichtbar zu machen und andererseits zu überprüfen, inwiefern diese Abweichungen durch die Interaktion im multikompetenten sprachlichen Repertoire (Cook, 1992, 2016) der Lernenden beeinflusst werden. Die orthographischen Abweichungen in den Lernertexten werden dabei im Sinne der Theorie dynamischer Systeme (CDST) als emergente Phänomene in einem komplexen, dynamischen System verstanden. Die quantitative und qualitative Analyse der Fehlerkategorien im gesamten DeChiLKo-Korpus hat ergeben, dass die orthographischen Schwierigkeiten chinesischer Deutschlernender nicht auf einzelne Bereiche beschränkt sind, sondern das gesamte Spektrum des deutschen Schriftsystems betreffen.

Eine Auswertung der absoluten Fehlerfrequenzen liefert ein detailliertes Bild der häufigsten Fehlerquellen. Die folgende Rangliste zeigt die zehn quantitativ bedeutsamsten Fehlertypen:

Tabelle 48: Die zehn häufigsten Fehlerarten und ihre relativen Fehlerwerte (pro 1.000 Tokens) im DeChiLKo-Korpus

Rang	Ebene	Fehlerart	Fehlerwert pro 1.000 Tokens
1	PHO	Ausgelassene Konsonantengrapheme	24,78
2	PHO	Falsche Konsonantengrapheme	22,95
3	PHO	Falsche Vokalgrapheme	22,95
4	SON	Grammatisch abzuleitende Deklinationsfehler	19,20
5	PHO	Überflüssige Konsonantengrapheme	18,09
6	SYN	Getrennt- für Zusammenschreibung	17,59
7	PHO	Falsches Reduktionsvokalgraphem	15,25
8	PHO	Ausgelassenes Reduktionsvokalgraphem	15,25
9	SYN	Klein- für Großschreibung bei Substantiven	10,17
10	PHO	Überflüssiges Reduktionsvokalgraphem	8,93

Die Analyse dieser Rangliste offenbart zwei zentrale Befunde. Erstens kristallisiert sich eine klare Dominanz von Abweichungen im phonographischen Bereich (PHO) heraus: Sieben der zehn häufigsten Fehlertypen entfallen auf diese Ebene. Hierzu gehören vor allem Fehler bei der Konsonantenschreibung sowie bei der Verschriftung von Vokalen, insbesondere der Reduktionsvokale. Zweitens wird jedoch deutlich, dass auch Phänomene der syntaktischen Ebene hochgradig fehleranfällig sind. Mit der Getrennt- statt Zusammenschreibung und der Kleinschreibung von Substantiven finden sich zwei syntaktische (SYN) Fehler unter den häufigsten Abweichungen. Diese Verteilung unterstreicht, dass die Herausforderungen für die Lernenden sowohl auf der grundlegenden Laut-Buchstaben-Ebene als auch auf der komplexeren syntaktischen Ebene angesiedelt sind.

Wie die Auswertung der Fehlerfrequenzen bereits andeutete, stellen die Abweichungen im phonographischen Bereich die größte Herausforderung dar. Im phonographischen Bereich zeigten sich besonders hohe Fehlerwerte bei der Phonem-Graphem-Korrespondenz und der Vokalquantitätsmarkierung. Die qualitative Analyse bestätigt eindrücklich die Annahmen der Multi-Kompetenz-Hypothese. Diese besagt, dass die verschiedenen Sprachen eines Lernenden (hier: Chinesisch, *Pinyin*, Englisch und Deutsch) im Gehirn nicht isoliert gespeichert sind, sondern ein einziges, vernetztes Gesamtsystem bilden (vgl. Cook, 2012, 2016). Dieses integrierte Sprachwissen agiert beim Erwerb einer weiteren Sprache sowohl als Ressource als auch als Quelle von Interferenz.

Einerseits erweist sich das Vorwissen aus *Pinyin* und Englisch als positive Resource. Da sowohl die *Pinyin*-Umschrift des Chinesischen als auch die englische Sprache das lateinische Alphabet verwenden, verfügen chinesische Lernende bereits bei Studienbeginn über eine Vertrautheit mit dem visuellen Erscheinungsbild und den motorischen Abläufen beim Schreiben der meisten deutschen Grapheme. Dieses vorhandene graphematische Wissen wirkt als starke positive Transfergrundlage und führt zu einer nahezu fehlerfreien Beherrschung der Buchstabenformen (vgl. Kapitel 4.2.1).

Andererseits ist das multi-kompetente Repertoire die plausible Ursache für systematische Interferenzfehler. Die Analyse der vorliegenden Daten hat hier drei zentrale Quellen offengelegt. Erstens ist, was auch im Einklang mit etablierter Forschung steht, die fehlende distinktive Vokallänge im L1-System des Chinesischen die Ursache für persistente Schwierigkeiten bei der deutschen Vokalquantitätsmarkierung (vgl. Kapitel 4.2.3). Mit diesem Ergebnis bestätigt die vorliegende Arbeit einen zentralen Befund der DaF-/DaZ-Rechtschreibforschung. Die Schwierigkeit bei der Markierung der Vokalquantität stellt eine typische Hürde für Lernende dar, deren Erstsprachen dieses Merkmal nicht kennen. Entsprechende Befunde liegen bereits für Lernende mit Türkisch (Becker, 2013; Steinig et al., 2009), Russisch (H.-G. Müller & Schroeder, 2022) oder Italienisch (Hans-Bianchi & Katelhön, 2010; Richter, 2008) als L1 vor. Die vorliegende Analyse liefert somit den ersten systematischen Beleg dieser typischen Lernhürde für chinesische Deutschlernende und stützt damit frühere Beobachtungen von T. Liu (2015), M. Liu (2019) und Tang (2003).

Eine zweite Interferenzquelle ist der Einfluss der Orthographiesysteme von *Pinyin* und der L2 Englisch. Dieser führt zu spezifischen Graphemverwechslungen (z. B. <z> statt <s> beim Laut [z], vgl. Kapitel 4.2.4) und differenziert die wissenschaftliche Debatte zur Rolle des L1-Transfers. Während einige Studien im DaZ-Kontext den direkten Einfluss erstsprachlicher Schreibkonventionen als eher gering einschätzen (vgl. u. a. Griefhaber, 2004), weisen andere, insbesondere im Kontext Chinesisch-Deutsch, explizit auf den negativen Transfer aus der L2 Englisch hin (vgl. Chi, 2016; Ding, 2018). Die vorliegende Studie beweist, dass der Einfluss hochspezifisch ist und sich nicht als generelle Interferenz, sondern in punktuellen Verwechslungen manifestiert, die auf konkrete Ähnlichkeiten oder Unterschiede im Grapheminventar von *Pinyin* oder Englisch im Vergleich zur Zielsprache Deutsch zurückzuführen sind.

Drittens greift die vorliegende Studie auf die Beobachtung von M. Liu (2019) zur Rolle der L1-Dialekte zurück und untermauert diese mit systematischen Daten. Die Analyse von Verwechslungen wie <n> für <ng> oder <l> für <n> (vgl. Kapitel 4.2.2) macht deutlich, dass die dialektale Prägung als Teil des multi-kompetenten Repertoires die Wortschreibung beeinflussen kann. Die vorliegende Untersuchung leistet hier einen Beitrag zur Erklärung der hohen interindividuellen Variabilität im Lernprozess, indem sie nachweist, wie spezifische dialektale Anfangsbedingungen zu unterschiedlichen Lernpfaden und Fehlerprofilen führen.

Zusammengenommen belegen die Funde, dass sich die beobachteten Fehlerprofile aus einem komplexen Zusammenspiel ergeben. Die hier diskutierten lernerseitigen Faktoren (positive und negative Transfereffekte aus L1, L1-Dialekt und L2) interagie-

ren dabei unmittelbar mit der inhärenten Komplexität des deutschen Orthographiesystems selbst. Regelungsbereiche wie die f-/v-Schreibung, die komplexe Systematik der s-Schreibung oder die oft nicht eindeutig aus der Lautung ableitbare Kennzeichnung der Vokalquantität stellen allgemein für DaF-Lernende eine Lernhürde dar und können die aus dem mehrsprachigen Repertoire stammenden Unsicherheiten zusätzlich verstärken.

Die Analyse der Abweichungen im silbischen Bereich offenbart eine klar positionsabhängige Fehlerverteilung: Am fehleranfälligsten erweist sich der Silbenkern, gefolgt vom Endrand. Dass der Anfangsrand hingegen die geringste Fehlerdichte aufweist, entspricht der linguistischen Erwartung, da dieser Position im Deutschen kaum komplexe silbische Regeln zugrunde liegen und sie weitgehend dem phonographischen Prinzip folgt.

Dennoch fällt auch hier eine Dominanz von Substitutionsfehlern auf. Dabei zeigt sich am Anfangsrand das kontraintuitive Muster, dass Fehler bei einfachen Rändern häufiger auftreten als bei komplexen (vgl. Kapitel 4.3.1). Dies lässt sich einerseits mit der höheren Frequenz einfacher Strukturen im Deutschen, andererseits mit einer gezielten Aufmerksamkeitslenkung auf die als schwierig antizipierten Cluster erklären. Die qualitativen Fehler deuten auf spezifische L1-Interferenzen hin, wie die dialektal bedingte Verwechslung von <n> und <l> oder die phonetisch motivierte Substitution von <r> durch <h>. Gleichzeitig bestätigt die erwartete Reduktion von Clustern wie <str> zu <st> die grundlegenden Schwierigkeiten mit für das Chinesische atypischen Konsonantenfolgen.

Der Silbenkern ist maßgeblich durch die Problematik der Schwa-Laute ([ə], [ɐ]) geprägt (vgl. Kapitel 4.3.2), für die es im Chinesischen keine direkte Entsprechung gibt. Dies provoziert zwei gegenläufige, L1-basierte Strategien: Einerseits die epenthetische Hinzufügung eines überflüssigen <e>, um die im Chinesischen präferierte offene Silbenstruktur zu erzeugen, andererseits die Auslassung des Schwas, insbesondere vor silbischen Sonoranten.

Während deutsche Diphthonge dank positiver Transfereffekte aus der vokalreichen L1 Chinesisch kaum Schwierigkeiten bereiten, ist die Schreibung von Vokalen mit nachfolgendem vokalisiertem <r> hochgradig fehleranfällig. Zudem zeichnet sich eine interessante Asymmetrie ab: Fehler bei den Langvokalen betreffen primär die Quantitätsmarkierung (z. B. <i> statt <ie>), während Fehler bei Kurzvokalen eher die Vokalqualität betreffen (z. B. <u> statt <ü>).

Der Endrand erweist sich als hochgradig fehleranfällig, was auf den Kontrast zwischen dem hochkomplexen deutschen Silbenauslaut und der stark eingeschränkten Struktur im Chinesischen zurückzuführen ist (vgl. Kapitel 4.3.3). Der dominanteste Transfereffekt ist das Hinzufügen eines überflüssigen <n>, eine direkte Übergeneralisierung des häufigsten chinesischen Auslautkonsonanten. Eine weitere L1-geprägte Fehlerquelle ist die häufige Auslassung von Sonoranten im Silbenendrand, insbesondere von <n>, <l> und <m>. Da im Chinesischen nach dem Hauptvokal nur die Nasale /n/ und /ŋ/ als Auslautkonsonanten erlaubt sind, werden andere im Deutschen mögliche Endkonsonanten von den Lernenden auditiv nur schwer perzipiert und daher

in der Schrift ausgelassen. Bei der Schreibung des Silbengelenks, also der korrekten Markierung der Vokalquantität durch Konsonantenverdopplung nach Kurzvokal, bestehen generelle Probleme. Dies betrifft sowohl deutsche Erbwörter (z. B. *<komen> statt *kommen*) als auch Entlehnungen aus dem Englischen (wie z. B. *<Hoby> statt *Hobby*).

Im morphologischen Bereich erweist sich die Affixschreibung als die quantitativ größte Fehlerquelle, während Verstöße gegen die Stammkonstanz auf die tiefgreifendsten konzeptuellen Herausforderungen hindeuten (vgl. Kapitel 4.4.1). Damit bestätigt sich für die chinesischen Lernenden im DaF-Kontext eindrucklich, was bereits für andere Lernergruppen im DaZ-Kontext als zentraler Problembereich identifiziert wurde (vgl. Fix, 2002; Ruppert & Hanulíková, 2022). Diese besondere Herausforderung der Stammkonstanz wird auch dadurch unterstrichen, dass sie selbst beim deutschen L1-Erwerb zu den größten Hindernissen zählt (Berg et al., 2023; Fix, 2002). Die systematischen Fehler bei der Umlautung und der Auslautverhärtung zeigen, dass die Lernenden oft auf eine rein phonetisch orientierte Verschriftung zurückgreifen. Dieser Rückgriff auf eine phonetische Strategie ist plausibel, da das abstrakt-orthographische Prinzip der Morphemkonstanz, welches eine Schreibung entgegen dem phonographischen Prozess wie der Auslautverhärtung fordert, im sprachlichen Repertoire der chinesischen Lernenden konzeptuell nicht verankert ist. Als Lernende einer isolierenden Sprache müssen sie dieses Prinzip der deutschen Orthographie als neues, abstraktes Regelwissen explizit aufbauen. Die dabei häufig beobachtete Instabilität, bei der derselbe Lernende einen Wortstamm inkonsistent verschriftet, ist jedoch kein exklusives Merkmal dieser Lernergruppe. Vielmehr handelt es sich um ein auch im Erstspracherwerb weit verbreitetes Phänomen, das die generelle Nicht-Linearität dieses Lernprozesses und die hohe kognitive Anforderung der Stammkonstanzschreibung zusätzlich unterstreicht.

Die konzeptuellen Lücken, die sich bereits bei der Stammkonstanz zeigten, prägen in ähnlicher Weise auch den Bereich der Wortbildung (vgl. Kapitel 4.4.2). Die enorme Fehleranfälligkeit bei Fugenelementen und Affixen steht in engem Zusammenhang mit der Tatsache, dass diese expliziten morphologischen Marker in der L1 Chinesisch fehlen. Diese Beobachtung deckt sich mit den Befunden anderer Studien zu chinesischen Deutschlernenden, die ebenfalls Fehler bei der Derivation (vgl. Ding, 2018) sowie bei Wortbildungselementen wie Affixen und Fugen deskriptiv erfasst haben (vgl. Z. Wu & Li, 2022). Diese L1-bedingte Schwierigkeit trifft dabei auf ein Zielsprachliches System, das selbst durch eine hohe Komplexität und Unregelmäßigkeit, etwa bei der Fugenvergabe, gekennzeichnet ist und den Erwerb zusätzlich erschwert. Die Abweichungen im morphologischen Bereich sind somit das emergente Ergebnis eines Restrukturierungsprozesses, der hohe kognitive Anforderungen an die Lernenden stellt. Sie müssen für die Zielsprache Deutsch ein neues morphologisches System mit Kategorien wie Flexion und Derivation entwickeln. Dabei ist methodisch zu berücksichtigen, dass diese Kompetenz im Diktat zwar primär über die Verschriftung akustischer Vorgaben erfasst wurde, die systematischen Abweichungen deuten jedoch darauf hin, dass diese Kategorien ko-

gnitiv noch nicht hinreichend verankert sind, was auch damit zusammenhängen dürfte, dass in ihrer isolierenden L1 keine konzeptuelle Entsprechung existiert.

Die Analyse der syntaktischen Ebenen belegt, dass die orthographischen Herausforderungen für die Lernenden vor allem in Regelungsbereichen liegen, die eine korrekte syntaktische Analyse des Satzkontexts erfordern. Die Getrennt- und Zusammenschreibung (GZS) stellt dabei die quantitativ größte Fehlerquelle dar, gefolgt von der Groß- und Kleinschreibung (GKS). Damit bestätigt die vorliegende Studie eindrücklich, „dass die syntaktische Schreibung zu den schwierigsten Bereichen der deutschen Orthographie für *alle* Schüler*innen gehört“ (Nimz, 2022, S. 47).

Dass sich die Getrennt- statt Zusammenschreibung als die dominante Fehlerquelle im syntaktischen Bereich erweist (vgl. Kapitel 4.5.2), unterstreicht die besondere Relevanz eines Problems, das auch im deutschen L1-Erwerb nachweisbar ist (vgl. Fuhrhop & Romstadt, 2021, S. 190). Für die chinesischen Lernenden wird diese inhärente Schwierigkeit des Deutschen jedoch durch L1-spezifische Faktoren auf einer visuell-konzeptuellen und einer auditiv-prosodischen Ebene zusätzlich verstärkt. Auf der visuellen Ebene ist das Wortverständnis der Lernenden durch ihre L1 geprägt, in der die visuelle Einheit eines Schriftzeichens oft einer semantisch-morphematischen Einheit entspricht. Auf der auditiven Ebene kommt hinzu, dass sich die Wortbetonungssysteme des Chinesischen und des Deutschen fundamental unterscheiden. Während im Deutschen sprachspezifische Betonungsmuster als entscheidendes Signal zur Abgrenzung von Wörtern im Sprachfluss dienen (vgl. Höhle, 2010, S. 130–131), ist das Chinesische eine Tonsprache. Dies bedeutet, dass die lexikalische Bedeutung primär durch den Tonhöhenverlauf einer Silbe und nicht durch den Wortakzent differenziert wird. Zwar existiert auch im Chinesischen eine Wortbetonung, doch ihre Funktion und Position sind flexibel. Bei isolierten Wörtern liegt die Betonung tendenziell auf der letzten tontragenden Silbe, doch in Sätzen und Wortgruppen dient sie vorrangig pragmatischen Zwecken, indem sie das sinnwichtigste Wort zur Fokussierung hervorhebt. Der Akzent ist im Chinesischen somit kein festes, zur Wortsegmentierung nutzbares Merkmal, sondern ein flexibles Werkzeug der Satzgestaltung (vgl. Hunold, 2021, S. 3).

Für chinesische Deutschlernende entsteht daraus eine doppelte Herausforderung: Weder ihr L1-geprägtes visuelles Wortkonzept noch ihre auditive Wahrnehmung bieten ihnen eine verlässliche Grundlage, um die Einheit eines komplexen deutschen Wortes zu erkennen. Die Befunde der vorliegenden Studie stützen somit die Annahme, dass der Einfluss der Erstsprache bei syntaktischen Fehlschreibungen hochrelevant ist. Diese Schwierigkeit in der orthographischen Umsetzung von Wortgrenzen wurde auch für andere mehrsprachige Lernergruppen nachgewiesen (vgl. Nimz, 2022, S. 46). Der L1-Einfluss erschwert somit auf mehreren Ebenen das Verständnis für mehrteilige Wortstrukturen, die im Deutschen als eine einzige orthographische Einheit verschriftet werden.

Eine weitere Herausforderung auf der syntaktischen Ebene sind die häufigen Verstöße gegen die Substantivgroßschreibung (vgl. Kapitel 4.5.1). Die Analyse legt nahe, dass die chinesischen Lernenden die morphosyntaktischen Signale, die im Deutschen eine Nominalisierung anzeigen (z. B. bei Abstrakta oder Substantivierungen wie *Wei-*

terlernen), noch nicht durchgängig als solche erkennen. Dieser Prozess wird durch das multikompetente Repertoire der Lernenden zusätzlich verkompliziert: Während die Vertrautheit mit der L2 Englisch die satzinitiale Großschreibung stützt, führt sie bei Nationalitätenadjektiven zu direkten Interferenzfehlern.

Hierbei zeigt sich ein entscheidender Unterschied zum deutschen L1-Erwerb. Während für Muttersprachler*innen die häufigste Fehlerquelle in diesem Bereich die Klein- statt Großschreibung ist (vgl. Fuhrhop & Romstadt, 2021; H. Günther & Gaebert, 2011; Röber, 2011, S. 296), weisen die chinesischen Lernenden zusätzlich eine signifikante Fehlerrate bei der Groß- statt Kleinschreibung auf, etwa bei deklinierten Adjektiven. Dies deutet darauf hin, dass die L3-Lernenden nicht nur Schwierigkeiten haben, die Regeln zur Substantivgroßschreibung anzuwenden, sondern aufgrund von Unsicherheit und Transferprozessen auch zur Übergeneralisierung neigen.

Die selteneren, aber qualitativ sehr aufschlussreichen Verwechslungen von Homophonpaaren wie *dass/das* und *Mann/man* (vgl. Kapitel 4.5.3) illustrieren eine grundlegende Verarbeitungsstrategie der Lernenden: Bei hoher kognitiver Last oder konzeptueller Unsicherheit greift das System auf eine basale, rein phonetische Verschriftung zurück, die als stabiler Attraktorzustand fungiert und die notwendige syntaktische Analyse überlagert. Diese lernerseitigen Tendenzen treffen zudem auf die inhärente Komplexität des Deutschen, etwa bei den fakultativen Regeln zur Bindestrichsetzung (insb. beim Ergänzungsstrich), die auch ohne weiteren Einfluss im Repertoire eine erhebliche Lernhürde darstellen.

Die Analyse der sonstigen Fehlerbereiche liefert einen der aufschlussreichsten Befunde der gesamten Untersuchung und unterstreicht die Notwendigkeit einer ganzheitlichen Betrachtung. Ein erheblicher Teil der orthographischen Abweichungen ist grammatisch bedingt. Dieser empirische Befund korrespondiert mit Eisenbergs (2020: 314) theoretischer Einordnung, dass Orthographiefehler oft „nichts anderes als ein spezieller Typ von Grammatikfehler“ sind, die darauf beruhen, „dass grammatische Eigenschaften eines Wortes nicht erkannt werden“. Diese theoretische Annahme spiegelt sich in der vorliegenden Studie konkret darin wider, dass nicht nur die explizit grammatischen Fehler ([Son_Gra]), sondern auch Abweichungen auf den anderen Ebenen (z. B. falsche Flexionsendungen) auf fehlendes grammatisches Wissen zurückzuführen sind.

Die besondere Anfälligkeit der deutschen Flexionsmorphologie ist dabei kein reines L2-Phänomen, wie Berg et al. (2023) belegen: Selbst bei deutschen Muttersprachler*innen stellen Flexionssuffixe, insbesondere die syntaktisch gesteuerten und semantisch wenig salienten Kasusmarker, eine signifikante Schwachstelle für Auslassungsfehler dar. Die vorliegende Studie knüpft an diesen Befund an und demonstriert, wie diese bei L1-Sprecher*innen nachgewiesene, inhärente Komplexität durch die typologische Distanz zur isolierenden L1 Chinesisch massiv potenziert wird. Das Fehlen eines morphologischen Flexionssystems im L1-geprägten Repertoire stellt eine Anfangsbedingung dar, die dazu führt, dass eine bereits für Muttersprachler*innen fehleranfällige Struktur für diese Lernergruppe zu einem der größten Hindernisse wird.

Wie aktiv die Lernenden dabei auf ihr gesamtes multikompetentes Repertoire zurückgreifen, illustriert die Analyse der Fremdwortschreibung (vgl. Kapitel 4.6.1) in be-

sonderem Maße. Je nach Herkunft des Wortes werden unterschiedliche Wissenssysteme aktiviert: Bei englischen Lehnwörtern kommt es zur Übertragung von L2-Regeln oder zur Bildung hybrider Formen (z. B. *<Hobbies>), während bei französischen Wörtern eine Anpassung an deutsche Aussprachemuster erfolgt und bei chinesischen Lehnwörtern die *Pinyin*-Orthographie als direkte Transferquelle dient. Die Lernenden konstruieren somit aktiv eine funktionierende Lernenden-Orthographie, die das komplexe Ringen ihres multikompetenten Systems mit den Regularitäten der deutschen Schriftsprache widerspiegelt.

Die als „sonstige“ klassifizierten Fehler, wie Wortabbrüche oder sinnverändernde Substitutionen, deuten zudem darauf hin, dass die orthographische Performanz stark von kontextuellen Faktoren wie Prüfungsstress und Zeitdruck beeinflusst wird, die zu Kompensationsstrategien führen (vgl. Kapitel 4.6.3).

In Bezug auf Forschungsfrage eins lässt sich zusammenfassend festhalten, dass sich die orthographischen Herausforderungen chinesischer DaF-Lernender über das gesamte deutsche Schriftsystem erstrecken und sich im Kern auf zwei Ebenen manifestieren: Einerseits auf einer strukturell-phonologischen Ebene, was sich in Problemen mit der Vokalquantität und dem Silbenbau offenbart. Andererseits stellt die abstrakt-konzeptuelle Ebene eine tiefgreifendere Hürde dar, da hier morphologische und syntaktische Prinzipien wie die Morphemkonstanz oder die Groß- und Zusammenschreibung Anwendung finden.

Diese Fehlermuster sind im Rahmen der Multi-Kompetenz-Theorie als direktes Ergebnis der Interaktion innerhalb des gesamten sprachlichen Repertoires der Lernenden zu interpretieren. Das Vorwissen aus der L1 Chinesisch und der L2 Englisch wirkt hierbei sowohl als Ressource als auch als Interferenzquelle. Die spezifischen Abweichungen sind somit keine reinen Transferphänomene, sondern emergieren aus dem dynamischen Zusammenspiel der L1/L2-geprägten Subsysteme auf unterschiedlichen Ebenen mit den inhärenten Komplexitäten der deutschen Orthographie.

6.1.2 Die Dynamik des Orthographieerwerbs

Zur Beantwortung von Forschungsfrage zwei wurde mithilfe der longitudinalen Daten des DeChiLKo-Erwerbskorpus und mittels statistischer Modellierungen untersucht, wie sich die orthographische Kompetenz chinesischer Deutschlernender in der Anfangsphase entwickelt. Die Analyse offenbart einen Prozess, der weit von einer einfachen, linearen Verbesserung entfernt ist. Stattdessen erweist er sich als hochgradig nicht linear, bereichsspezifisch und individuell variabel.

Im ersten Semester dominieren grundlegende phonographische (PHO) und silbische (SIL) Fehler. Ab dem zweiten Semester verlagern sich die Hauptschwierigkeiten auf komplexere syntaktische (SYN) und insbesondere grammatiknahe „sonstige“ Fehler (SON), deren Anteil signifikant zunimmt. Dieses im zweiten Semester etablierte Fehlerprofil bleibt in der Folgezeit weitgehend stabil. Der Lernfortschritt verläuft zudem gestaffelt und asynchron, da signifikante Verbesserungen im silbischen Bereich früh, im phonographischen und morphologischen Bereich jedoch erst später einset-

zen, während in den komplexesten syntaktischen und grammatiknahen Bereichen über den Beobachtungszeitraum kaum Fortschritte zu verzeichnen sind.

Diese Entwicklungsreihenfolge stimmt jedoch nur sehr eingeschränkt mit den Ergebnissen bestehender Studien zum L1-Schriftspracherwerb überein. Bulut (2018, S. 109) sowie A. Reichardt (2015, S. 202–203) belegten in ihren Studien für den L1-Orthographieerwerb die Reihenfolge: phonographisches Prinzip, morphologisches Prinzip und schließlich silbisches Prinzip, wobei Kinder die meisten Probleme mit der silbischen Schärfungsschreibung haben. Der Befund aus dem Erwerbskorpus, dass sich die silbischen Kompetenzen der Lernenden verhältnismäßig früh verbessern, stellt einen markanten Unterschied zur L1-Entwicklung dar. Dieser Befund deutet stark darauf hin, dass die Anfangsbedingungen der mehrsprachigen, erwachsenen Lernenden fundamental andere sind. Faktoren wie eine höhere kognitive Reife, ein explizit regelbasierter DaF-Unterricht sowie ein durch *Pinyin*- und Englisch-Vorkenntnisse geschärftes silbisches Bewusstsein könnten den Erwerb dieser Prinzipien im Vergleich zum impliziten L1-Erwerb von Kindern beschleunigen. Die von den Lernenden im syntaktischen Bereich gezeigten Schwierigkeiten sind ein breit belegtes Phänomen, das nicht auf L3-Lernende beschränkt ist. So hat beispielsweise Nimz (2022, S. 43) statistisch belegt, dass die syntaktisch relevanten Fehlerkategorien in allen beobachteten Klassen (5, 7 und 10 Klassen) sehr hohe Fehlerquotienten aufweisen und signifikante Unterschiede zwischen bilingualen und monolingualen Schüler*innen zeigen. Auch Steinig et al. (2009) haben eine höhere Fehlerquote für Kinder mit Deutsch als L2 bei der Klein- statt Großschreibung festgestellt.

Die beobachtete Entwicklungsdynamik, insbesondere die von der L1-Forschung abweichende Reihenfolge, stützt die Annahmen der CDST. Der Unterschied in der Erwerbsreihenfolge deutet darauf hin, dass die Emergenz der Zustandsräume eines mehrsprachigen Individuums anders verläuft als die eines einsprachigen Kindes (vgl. Bredel, 2013, S. 376). Die Herausbildung dieser neuen, individuellen Zustandsräume entsteht aus der dynamischen Interaktion der spezifischen Komponenten im multikompetenten Repertoire, die zudem von externen Faktoren beeinflusst wird.

Trotz dieser allgemeinen Gruppentendenzen wird eine hohe interindividuelle Variabilität deutlich, da die Entwicklungsverläufe der einzelnen Probanden deutlich voneinander abweichen und individuelle Profile wie Fehlergipfel im zweiten Semester oder ausgeprägte „Umkehr-U-Kurven“ aufweisen. Diese Beobachtungen unterstreichen zudem, dass die Zeitpunkte der größten Dynamik stark variieren und auch die dominanten Fehlertypen lernerspezifisch sind. Während bei einigen Lernenden (LWY, YYYY) Fehler im sonstigen Bereich (SON) dominieren, springen bei anderen syntaktische (ZXY) oder phonographische (ZZX) Fehler hervor. Auch gibt es individuell „stille“ Bereiche, wie den silbischen Bereich (SIL) bei LWY, der sich kaum verändert.

In Bezug auf Forschungsfrage zwei lässt sich zusammenfassen, dass die orthographische Entwicklung der DaF-Lernenden nicht linear verläuft, sondern ein komplexer, dynamischer Prozess ist. Es ergibt sich eine klare Verschiebung der Fehlerprofile von phonographisch-silbischen Aspekten hin zu komplexeren syntaktischen und grammatischen Herausforderungen. Zudem zeigt sich eine gestaffelte, asynchrone Verbesse-

rung der einzelnen Teilkompetenzen. Einer der wichtigsten Befunde ist, dass der hier beobachtete Erwerbspfad, insbesondere die Entwicklungsreihenfolge, von den Modellen des L1-Orthographieerwerbs abweicht. Die Einzelfallanalysen der fünf Proband*innen im Erwerbskorpus bewiesen, dass sich ihre beobachtbare lernersprachliche Rechtschreibentwicklung stark voneinander unterscheidet. Diese stark unterschiedlichen Entwicklungsverläufe sind eine direkte Konsequenz der Tatsache, dass jede*r Lerner*in von qualitativ anderen und einzigartigen Anfangsbedingungen ausgeht. Das individuelle multikompetente Repertoire und die persönliche kognitive Reife bedingen somit bei jedem Einzelnen einen eigenständigen Entwicklungspfad.

6.1.3 Das Zusammenspiel der Einflussfaktoren

Die Forschungsfrage drei untersuchte, welche Faktoren die orthographische Leistung chinesischer DaF-Lernender beeinflussen und wie deren dynamisches Zusammenspiel zu charakterisieren ist. Zur Beantwortung dieser Frage wurde auf die Daten beider DeChiLko-Subkorpora zurückgegriffen und eine differenzierte statistische Analyse mittels eines zweistufigen Hurdle-Modells sowie Likelihood-Quotienten-Tests durchgeführt. Dieses Verfahren ermöglichte es, zwischen der reinen Fehlerwahrscheinlichkeit und der tatsächlichen Fehlerhäufigkeit zu differenzieren und so die Wirkung der einzelnen Faktoren präziser zu bestimmen.

Im Zentrum des Modells für das DeChiLko-Erwerbskorpus steht der Lernende selbst und sein individueller Entwicklungsprozess. Die Längsschnittanalyse des Erwerbskorpus liefert hierfür den entscheidenden Beleg: Der Faktor *Lerndauer* erweist sich als hochgradig signifikanter Einflussfaktor auf die orthographische Leistung. Dieses statistische Ergebnis bestätigt zunächst die grundlegende Annahme aus der Perspektive der CDST, dass der orthographische Erwerb ein zeitabhängiger Prozess ist (vgl. de Bot & Larsen-Freeman, 2011, S. 11), in dem sich die Kompetenz der Lernenden verändert und nicht statisch bleibt. Wie bereits in Abschnitt 6.1.2 detailliert dargelegt wurde, bedeutet dies keineswegs eine einfache, lineare Verbesserung. Die Entwicklung ist vielmehr durch komplexe, nicht lineare Dynamiken wie gestaffelte Fortschritte in unterschiedlichen Bereichen, temporäre Regressionen und sich verschiebende Fehlerschwerpunkte gekennzeichnet. Die Ergebnisse müssen vor dem Hintergrund einer „curricularen Leerstelle“ betrachtet werden: In den untersuchten Germanistikstudiengängen an der Anhui-Universität fehlt sowohl systematischer Orthographieunterricht als auch die Vermittlung von Orthographieprinzipien (vgl. Kapitel 5.1.4). Dies deutet darauf hin, dass die beobachtete, nicht lineare Kompetenzentwicklung primär durch einen Prozess der Selbstorganisation im kognitiven System der Lernenden gesteuert wird.

Die Analyse des Erwerbskorpus zum Faktor *Aufgabentyp* liefert zudem einen weiteren Einblick in die Natur der orthographischen Kompetenz selbst. Die Tatsache, dass situativer Prüfungsstress die orthographische Leistung nicht signifikant beeinflusst, steht im deutlichen Kontrast zu den nachgewiesenen negativen Effekten von Angst auf komplexere sprachliche Aufgaben. So belegen Studien, dass Schreibangst zu qualitativ schlechteren Texten mit mehr Fehlern führt, da kognitive Ressourcen durch die Angst gebunden werden (Rasool et al., 2023; Zheng & Cheng, 2018). Eine entscheidende Dis-

soziation zeigt sich jedoch bei grundlegenden, stärker automatisierten Fertigkeiten. Dies bestätigt eine Studie von Kargiotidis und Manolitsis (2024), die keinen statistisch signifikanten Zusammenhang zwischen allgemeiner Testangst und basalen Fertigkeiten wie der Worterkennungsgenauigkeit fand, während gleichzeitig komplexere Leistungen wie das Textverständnis beeinträchtigt wurden. Das Ergebnis der vorliegenden Untersuchung deckt sich mit diesen Befunden: Die orthographische Performanz erweist sich hier, ähnlich der basalen Worterkennung, als eine relativ stabile Kompetenz, die von situativem Stress wenig beeinflusst wird.

Im Sinne der CDST bedeutet dies, dass sich das System in einem Zustand befindet, in dem die typischen Performanzmuster bereits gut etablierte Attraktoren sind (vgl. de Bot et al., 2007; Hiver, 2014; Wind & Harding, 2020). Diese Interpretation wird durch Forschungen wie die von Wind und Harding (2020) gestützt, die eine hohe Stabilität (geringe Variabilität) in der linguistischen Komplexität von Texten als direkten Hinweis auf das Vorhandensein von „markanten Attraktorzuständen“ (*salient attractor states*) werteten. Eine situative Störung wie Prüfungsdruck genügt somit nicht, um das System aus diesen stabilen Zuständen abzulenken (vgl. de Bot et al., 2007; Larsen-Freeman, 2017). Dem Stressor fehlt die „nachhaltige oder energische Kraft“ (*sustained or vigorous force*) (Hiver, 2014, S. 25), die notwendig wäre, um das System aus seinem tiefen Attraktorbecken zu stoßen und eine signifikante Veränderung in der Gesamtperformanz hervorzurufen.

Der individuelle Erwerb orthographischer Kompetenz findet jedoch nicht im luftleeren Raum statt, sondern ist in institutionelle und regionale Kontexte eingebettet, die als mächtige Rahmenbedingungen fungieren. Diese Beobachtung deckt sich mit umfangreichen Forschungen zum chinesischen Bildungssystem, die ein signifikantes strukturelles Ungleichgewicht aufzeigen. Bildungsressourcen und hochrangige Institutionen sind stark in den östlichen Metropolen wie Peking und Shanghai konzentriert, während die zentralen und westlichen Regionen deutlich zurückfallen (vgl. Goldberger, 2017; Q. Wu & Borhan, 2024). Die Querschnittsanalyse des DeChiLKo-Prüfungskorpus bestätigte diese strukturellen Realitäten: Hochschulkategorie und geographische Region beeinflussen signifikant sowohl die Fehlerwahrscheinlichkeit als auch die Fehlerhäufigkeit im Bereich der Orthographie (vgl. Kapitel 5.1.3 und Kapitel 5.1.5).

Obwohl es bisher noch keine Untersuchung gibt, die die Korrelation des *Hochschultyps* und der Orthographiekompetenz von DaF-Lernenden untersucht, wurde der Einfluss des Schultyps auf die Rechtschreibkompetenz bereits durch zahlreiche Studien belegt (vgl. u. a. Fix, 2002; H.-G. Müller & Schroeder, 2022). Große Vergleichsstudien wie der IQB-Bildungstrend 2022 belegen, dass die Schulart den statistisch stärksten Einfluss auf die Orthographiekompetenz hat. Dieser Faktor erweist sich als bedeutsamer als andere interne Merkmale wie kognitive Grundfähigkeiten oder Zuwanderungshintergrund sowie externe Faktoren wie die Unterrichtsgestaltung (vgl. Henschel, Rjosk & Heinschel, 2023, S. 376). Es muss berücksichtigt werden, dass die Schulart in Deutschland als ein starker Indikator für den sozioökonomischen Status (SES) gilt (vgl. Bulut, 2018, S. 110). Großangelegte Studien wie die PISA-Erhebungen (vgl. Maaz & Lörz, 2024) und IQB-Berichte (vgl. Stanat et al., 2023) belegen konsistent die enge Korrelation zwi-

schen der besuchten Schulform und sozialen Hintergrundmerkmalen der Schülerschaft. So weisen Studien von Voss, Blatt und Kowalski (2007) und Niemietz et al. (2023) nach, dass das Leseangebot im Elternhaus positiv mit der Rechtschreibkompetenz im Erstspracherwerb korreliert. Da die vorliegende Studie den Fokus jedoch nicht auf die Erhebung solcher spezifischer SES-Variablen – beispielsweise den Bildungshintergrund der Eltern oder Lesegewohnheiten – legte, können diese hier nicht vergleichend diskutiert werden.

Auf Basis dieser Befunde lassen sich die Faktoren *Hochschulkategorie* und *Region* als Indikatoren für zwei zentrale Aspekte interpretieren: die allgemeinen Lernvoraussetzungen der Studierenden, die durch die Gaokao-Zulassung vorselektiert werden, und die Qualität der jeweiligen Lernumgebung. Die Interpretation der Gaokao als Indikator für Lernvoraussetzungen korrespondiert mit deren gesellschaftlicher Wahrnehmung in China, wo der Prüfungserfolg als entscheidendes Kriterium für die „Lernfähigkeit“ und Intelligenz von Bewerbern gilt (vgl. Ren, 2022). Dass die kognitive Leistungsfähigkeit einen signifikanten Einfluss auf die Rechtschreibkompetenz hat, wurde für Schülerinnen und Schüler bereits mehrfach empirisch nachgewiesen (Becker, 2013; Corvacho Del Toro, 2013; Henschel et al., 2023; z. B. Roos & Schöler, 2009). Auch die Längsschnittstudie von Bulut (2018) stützt diese Sichtweise, indem sie empirisch zeigt, wie eine höhere kognitive Grundfähigkeit mit einer früheren Bewältigung von Rechtschreibphänomenen einhergeht. Diese Befunde bekräftigen das CDST-Prinzip der Selbstorganisation, das an individuelle kognitive Voraussetzungen geknüpft ist und somit erklärt, warum manche Lernende bestimmte orthographische Regeln früher als andere beherrschen (vgl. Bulut, 2018, S. 116).

Die beiden Faktoren *Hochschulkategorie* und *Region* markieren somit die qualitativen externen Anfangsbedingungen im Sinne der CDST, unter denen der individuelle Lernprozess stattfindet. Studierende an Elite-Universitäten in ressourcenstarken Regionen starten folglich nicht nur mit tendenziell besseren kognitiven Voraussetzungen, sondern profitieren auch von einer Umgebung, die den selbstorganisierten Lernprozess effektiver unterstützt.

Schließlich erlauben die Analysen der Forschungsfrage 3 eine Differenzierung von situativen und strukturellen Einflüssen. Der Faktor *Prüfungsstatus* erweist sich im DeChiLKo-Prüfungskorpus als hochgradig signifikanter Prädiktor sowohl für die Wahrscheinlichkeit des Fehlerauftretens als auch für die Höhe der Fehlerwerte. Die beobachtete Variabilität der Fehlerwerte darf dabei jedoch nicht als Indikator für eine tatsächliche Entwicklungsregression oder einen Rückgang der Kompetenz missverstanden werden. Vielmehr deckt die Analyse hier einen methodischen Selektionseffekt auf: Bei den Lernenden der späteren Studienjahre handelt es sich primär um Prüfungswiederholer*innen, die trotz längerer Vorbereitungszeit systematisch höhere orthographische Fehlerwerte aufweisen als der reguläre Jahrgang. Auch der Faktor *Diktattext* entfaltet auf beiden Stufen der Fehlerproduktion eine maßgebliche Wirkung (vgl. Kapitel 5.1.2). Ein schwierigerer Text erhöht somit nicht nur die Chance, dass ein orthographischer Fehler überhaupt gemacht wird, sondern determiniert durch seine linguistischen Schwer-

punkte massiv, in welchen Bereichen und in welchem Ausmaß orthographische Fehler insgesamt produziert werden.

Zusammenfassend lässt sich anhand der Frage 3 feststellen, dass der L3-Orthographieerwerb chinesischer Deutschlernender als emergentes Phänomen zu verstehen ist, das aus dem komplexen Zusammenspiel multipler Faktoren resultiert. Im Kern dieses Prozesses steht die individuelle Entwicklungsdynamik. Diese wird maßgeblich von den externen, institutionellen und regionalen Anfangsbedingungen moderiert, wobei zu beachten ist, dass diese externen Faktoren nicht vollständig von den individuellen Voraussetzungen der Lernenden zu trennen sind. Insbesondere der Faktor der *Hochschulkategorie* ist bereits eng mit der kognitiven Leistungsfähigkeit der Studierenden verwoben. Aus dieser dynamischen Interaktion zwischen dem sich selbst organisierenden Lerner*innen und der Qualität ihrer Lernumgebung bildet sich schließlich eine relativ stabile orthographische Kompetenz zu jedem Messzeitpunkt heraus, die von rein situativen Performanzschwankungen, etwa durch Prüfungsstress, klar abzugrenzen ist.

Die vorliegende Arbeit hat auf Basis der drei zentralen Forschungsfragen dargelegt, dass der L3-Orthographieerwerb chinesischer Deutschlernender als ein komplexes dynamisches System zu verstehen ist. Somit entstand das konzeptuelle Modell in Abbildung 44, die das dynamische Zusammenspiel verschiedener Komponenten über die Zeit veranschaulicht.

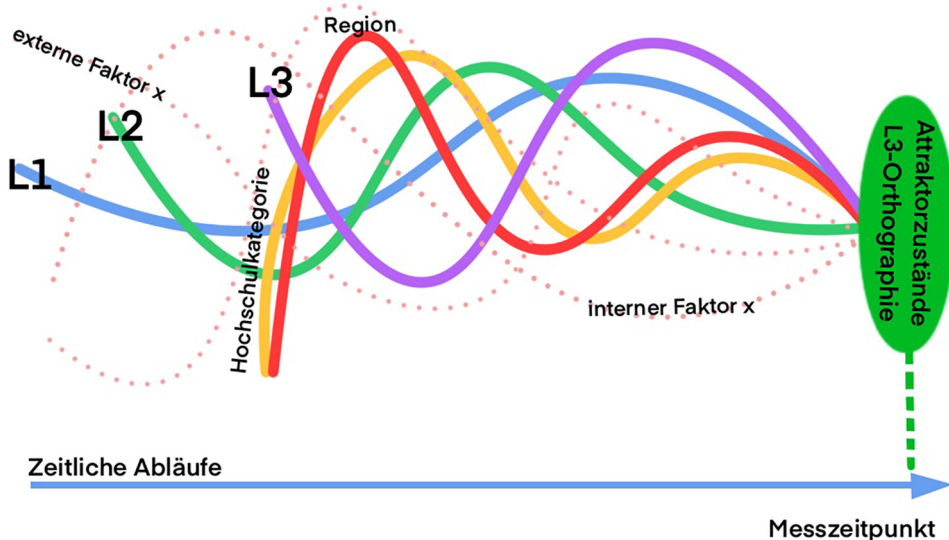


Abbildung 44: Konzeptuelles Modell des L3-Orthographieerwerbs chinesischer Deutschlernender

Auf der horizontalen Achse ist die zeitliche Entwicklung des Lernprozesses dargestellt. Die miteinander verwobenen, oszillierenden Linien (L1, L2, L3) repräsentieren das multikompetente Repertoire der Lernenden, in dem die verschiedenen sprachlichen Subsysteme kontinuierlich interagieren und sich gegenseitig beeinflussen. Die wellenförmige Bewegung symbolisiert die nicht lineare Entwicklungsdynamik; die vertikale

Position der Linien hat dabei keine quantitative Bedeutung, sondern veranschaulicht, dass der Lernprozess von konstanter Veränderung, Interaktion und Phasen der Restrukturierung geprägt ist, anstatt einem geradlinigen Pfad zu folgen.

In diesen Entwicklungsprozess wirken verschiedene Einflussfaktoren hinein. Externe Faktoren, die in dieser Arbeit untersucht wurden, wie die Region oder die Hochschulkategorie, fungieren als übergeordnete Rahmenbedingungen oder Anfangsbedingungen, die den Entwicklungsverlauf von außen moderieren. Die gepunkteten Linien deuten auf weitere, in dieser Studie nicht erfasste Faktoren (wie z. B. Motivation oder familiärer Hintergrund) hin, die ebenfalls auf das System einwirken und zu seiner Komplexität beitragen.

Über die Zeit hinweg tendiert das gesamte System in Richtung eines Attraktorbeckens. Im Laufe dieser Entwicklung kann das System in Phasen relativer Stabilität eintreten, sogenannte Attraktorzustände, die als dynamische Entsprechung zu den „Stufen“ klassischer Erwerbsmodelle verstanden werden können. Die an einem spezifischen Messzeitpunkt beobachtbare Performanz ist somit ein Abbild des aktuellen Zustands des Systems auf seinem Weg zu oder innerhalb dieser Attraktorzustände. Das Modell verdeutlicht somit, dass orthographische Kompetenz als eine emergente Eigenschaft aus dem komplexen Zusammenspiel all dieser Faktoren entsteht. Dabei ist zu betonen, dass jedes mehrsprachige Individuum einem eigenen, individuellen Entwicklungspfad folgt.

6.2 Didaktische Implikationen

Aus der Perspektive der Theorie komplexer dynamischer Systeme (CDST) und der Multi-Kompetenz-Hypothese wurde in der vorliegenden Arbeit deutlich, dass der Orthographieerwerb chinesischer Deutschlernender ein nicht linearer, hochgradig individueller und von der Interaktion multipler Wissenssysteme geprägter Prozess ist. Diese Erkenntnisse widersprechen einer simplen, defizitorientierten Fehlerkorrektur und erfordern eine prozessorientierte und diagnostisch fundierte Orthographiedidaktik. Dieses Leitprinzip lässt sich in konkrete Anforderungen an die Kompetenzen der Lehrenden und an die Gestaltung der Lerninhalte übersetzen.

Die Analyse im gesamten DeChiLKo-Korpus hat verdeutlicht, dass die orthographischen Abweichungen der chinesischen Lernenden keine zufälligen Flüchtigkeitsfehler sind, sondern systematische Phänomene mit multikausalen Ursachen. Eine Fehlschreibung wie z. B. *<füreisammen > für *vereinsamen* ist nicht nur ein einfacher Flüchtigkeitsfehler, sondern verweist gleichzeitig auf Probleme im phonographischen Bereich (Auslassung des Konsonantengraphems <n>; f/v-Schreibung), der Silbengelenkschreibung, der Präfixschreibung im morphologischen Bereich (*ver-*). Solche orthographischen Fehler sind das Ergebnis einer multikausalen Wechselwirkung: Einerseits führen strukturelle Unterschiede im sprachlichen Repertoire (L1 Chinesisch, L2 Englisch) zu Interferenzphänomenen, andererseits treffen diese auf die inhärente Komplexität des Deutschen. Orthographieerwerb unter den Bedingungen der Mehrsprachigkeit setzt

bereits Wissen über die Orthographien und Schriftsysteme von Sprachen im gesamten Repertoire voraus. Eine moderne Orthographiedidaktik sollte dieses Repertoire nicht ignorieren, sondern als Ressource für den mehrsprachigen Schriftspracherwerb nutzbar machen (vgl. Fuhrhop, 2022, S. 15). Wo systematische Fehler unter Einfluss von L1 oder L2 auftreten (z. B. bei der Getrennt- und Zusammenschreibung, Vokalquantität, GKS bei Nationalitätenadjektiven), ist ein expliziter kontrastiver Vergleich sinnvoll. Der Vergleich sollte sich jedoch nicht auf isolierte Phonem-Graphem-Beziehungen beschränken, sondern auch phonotaktische oder morphologische Strukturen einbeziehen (vgl. Fuhrhop, 2022, 18 f.). Fuhrhop (2022) illustrierte anhand der vergleichenden Graphematik im Englischen, Französischen, Niederländischen und Deutschen, dass ein und dasselbe Formmittel (wie z. B. das Trema) je nach sprachspezifischer Kombinatorik verschiedene Funktionen übernehmen kann. Das DeChiLKo-Korpus kann im DaF-Unterricht für chinesische Lernende zur Bewusstseinsbildung eingesetzt werden. Lehrende können damit kontrastive Analysen anleiten, in denen die Lernenden ihr eigenes sprachliches Repertoire reflektieren. So lässt sich beispielsweise durch gezielte Suchen zur Verwechslung von <n> und <ng> die Relevanz von L1-Dialekteinflüssen für die Lernenden sichtbar machen. Genau solche sprachvergleichenden „Entdeckungen“ der Systemhaftigkeit des sprachlichen Wissens im Repertoire der Lernenden können für den Orthographieerwerb im „mehrsprachigen Klassenzimmer“ (Krifka et al., 2014) eine Ressource sein (vgl. K. Schmidt, 2022, S. 8).

Eine pauschale Korrektur der Rechtschreibfehler, wie sie Lernende üblicherweise als Feedback zu einer Diktatübung im DaF-Unterricht bekommen, greift deshalb zu kurz. Stattdessen ist eine linguistisch fundierte Fehleranalyse durch die Lehrkräfte unerlässlich. Nur wenn eine Lehrkraft die Multikausalität eines Fehlers erkennt – ob er phonographisch, morphologisch, syntaktisch und/oder transferbedingt ist –, kann sie gezielte Fördermaßnahmen einleiten. Um diese diagnostische Kompetenz aufzubauen, bietet das in dieser Arbeit erstellte DeChiLKo-Korpus eine direkte Anwendungsmöglichkeit in der Lehrkräftefortbildung. Durch die Analyse von Beispielen aus dem Korpus können Lehrende die mehrschichtige Annotation eines einzigen Rechtschreibfehlers nachvollziehen. Gerade die Überlappungen zwischen den verschiedenen Annotationsebenen – wie bei der Fehlschreibung *<farat> für *Fahrrad* (siehe Abbildung 8) – machen die Multikausalität orthographischer Abweichungen sichtbar. Darüber hinaus kann das in dieser Studie entwickelte Annotationsschema von den Lehrkräften direkt in der eigenen Praxis angewendet werden, um die orthographischen Leistungen ihrer Schüler*innen systematisch zu erfassen, kleinere, klasseninterne Lernerkorpora aufzubauen und so die individuelle Entwicklung evidenzbasiert zu verfolgen und gezielte Fördermaßnahmen einzuleiten. Wie Bulut (2018) für den L1-Kontext darlegt, setzt effizientes Lernen voraus, dass der individuelle Kompetenzstand präzise erfasst wird.

Um das individuelle Rechtschreibwissen der Lernenden für die Lehrkraft sichtbar zu machen, bietet sich die Einführung des Rechtschreibgesprächs (RSG) an (vgl. Geist & Hesselbarth, 2017). Obwohl dieses Diagnoseinstrument ursprünglich für den L1-Orthographieunterricht entwickelt wurde, liegt hier großes Potenzial für den DaF-Kontext. Anstatt rein ergebnisorientierter Übungsformate wie Diktate bietet RSG

nicht nur einen Einblick in das (schrift-)sprachliche Wissen der Lernenden (vgl. Harnisch, 2014), sondern gibt auch Hinweise darauf, für welche Fehler die Lernenden (bereits) ein Bewusstsein haben (vgl. Geist & Hesselbarth, 2017, S. 146). Wie ein RSG in der DaF-Praxis mit chinesischen Lernenden konkret verlaufen kann, soll das folgende fiktive Beispiel anhand der Fehlschreibung *<füreisammen> für vereinsamen verdeutlichen:

KTN 1: „Ich habe *vereinsamen* geschrieben. Das Wort klingt für mich wie *ver-ein-sa-men* – da ist eine Betonung in der Mitte, bei *-ein-*.“

KTN 2: „Ich auch. Aber warum schreibt man das zusammen? *ver-* klingt doch fast wie ein eigenes Wort.“

KTN 1: „Ich glaube, *ver-* kann man nicht alleine sagen. Was bedeutet *ver* denn?“

KTN 3: „Nichts. Es funktioniert nur mit dem Verb zusammen – wie *verstehen*, *vergessen*, *verlieren*. Immer zusammen.“

KTN 2: „Ah, stimmt. Und das ist dann wie ein Präfix – das gehört fest dazu, auch wenn man es nicht betont.“

KTN 1: „Dann ist *füreisammen* falsch, weil für ein echtes Wort ist – das kann man alleine sagen. Aber *ver-* nicht.“

LK: (nickt, schreibt *vereinsamen* an die Tafel) „Was ist eure Regel?“

KTN 3: „Wenn man den Teil vor dem Verb nicht alleine sagen kann – dann schreibt man alles zusammen.“

Dieser Gesprächsausschnitt illustriert das Kernprinzip des RSG: Die Lernenden formulieren Hypothesen („*ver-* klingt fast wie ein eigenes Wort“), widerlegen und korrigieren sich gegenseitig und gelangen schließlich selbst zur Regel. Die Lehrperson greift erst am Ende ein – nicht um die Antwort zu geben, sondern um die Regelformulierung einzufordern (vgl. Geist & Hesselbarth, 2017, S. 146). Dieses kooperative Entdecken verbindet die prosodische Wahrnehmung – die Identifikation der Hauptsilbe und ihrer Umgebung – mit der morphologischen Analyse (Präfixerkennung). RSG hat sich in der DaF-Praxis bisher kaum etabliert. Ein wesentlicher Grund hierfür ist, dass der erfolgreiche Einsatz des RSG umfangreiches orthographisches Wissen vonseiten der Lehrkraft voraussetzt (vgl. Jagemann, 2016; Wiprächtiger-Geppert, Riegler & Freivogel, 2015), was wiederum die Notwendigkeit einer fundierten Lehrkraftaus- und Fortbildung unterstreicht (vgl. Jagemann, 2015).

Die Ergebnisse dieser Studie fordern eine Neubewertung der Lerninhalte im DaF-Unterricht. Anstatt eines reinen Fokus auf die Laut-Buchstaben-Ebene müssen die konzeptuellen Grundlagen der deutschen Orthographie in den Mittelpunkt rücken. Ein wichtiger Lerninhalt ist die Vermittlung der prosodischen Grundlagen der Orthographie. Wie die Fehleranalyse im Korpus belegt, lässt sich die Wichtigkeit hierfür direkt

aus der quantitativen Fehlerverteilung ableiten: Zentrale Schwierigkeiten (Reduktionsvokale, Vokalquantität, Getrennt- und Zusammenschreibung) sind auf eine mangelnde Wahrnehmung der rhythmisch-prosodischen Muster des Deutschen zurückzuführen. Da das Chinesische als Tonsprache den Wortakzent anders nutzt, entwickeln chinesische Lernende oft eine Form von „Betonungstaubheit“ (Bredel, 2013, S. 379). Sie können die für die deutsche Rechtschreibung entscheidenden Betonungsmuster nur schwer wahrnehmen und somit die darauf basierenden Regeln nicht intuitiv anwenden.

Eine linguistisch fundierte Didaktik muss daher, wie Bredel (2013) aufzeigt, über die reine Laut-Buchstaben-Ebene hinausgehen und die prosodischen Grundlagen der Orthographie explizit vermitteln. Für Lernende mit chinesischem Hintergrund bedeutet dies konkret, den für den deutschen Kernwortschatz typischen trochäischen Rhythmus (einer rhythmischen Abfolge aus einer betonten und einer unbetonten Silbe) des Deutschen zu thematisieren und vor allem zu visualisieren. Obwohl Methoden wie das „Silbenhaus“, das die betonte Hauptsilbe und die unbetonte Reduktionssilbe bildlich darstellt (vgl. Röber, 2013; Röber-Siekmeyer, 1997), oder die Kennzeichnung von Betonungsmustern durch Silbenbögen (vgl. Hinney, 1997) primär für Deutsch als L1 entwickelt wurden, können diese Modelle für die Lernenden mit Deutsch als Fremdsprache besonders lernförderlich sein, damit sie das trochäische Grundmuster „sehen“ und die typischen Wortstrukturen des Deutschen selbst entdecken und kognitiv verstetigen können (vgl. Bredel, 2013, 382 f.).

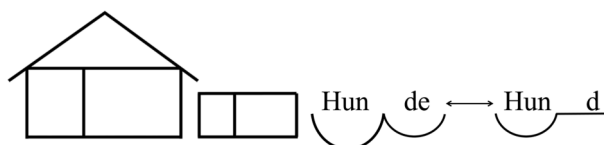


Abbildung 45: links: Didaktisches Basismuster nach Röber (2013) und rechts: Visualisierung von Silben-/Fuß- und Wortstrukturen nach Hinney (1997) (Quelle: Hinney, 1997; Röber, 2013)

Obwohl solche prosodiebasierten didaktischen Methoden bereits vor Jahrzehnten entwickelt wurden, finden sie in aktuellen DaF-Lehrwerken und im Unterrichtsalltag in China bislang kaum Anwendung. Die folgende Tabelle zeigt vier Fehlschreibungen aus dem DeChiLKo-Korpus, ihre Zielformen, die zugrunde liegende trochäische Struktur sowie den jeweiligen didaktischen Fokus. Sie kann direkt als Unterrichtsmaterial oder als Grundlage für eigene Silbenhaus-Aufgaben eingesetzt werden:

Tabelle 49: Ausgewählte Fehlschreibungen aus dem DeChiLKo-Korpus mit Silbenstruktur und didaktischem Hinweis

Fehlschreibung	Zielform	Silbenstruktur (Trochäus)	Didaktischer Hinweis
*Freizeit Angebote	Freizeitangebote	'FREI-zeit-an-ge-bo-te	Hauptsilbe liegt auf dem ersten Glied → kein Wortgrenzenzeichen
*Weiter Lernen	Weiterlernen	'WEI-ter-ler-nen	Trochäisches Muster zeigt Einheitlichkeit des Kompositums

(Fortsetzung Tabelle 49)

Fehlschreibung	Zielform	Silbenstruktur (Trochäus)	Didaktischer Hinweis
* <i>Sprache schule</i>	<i>Sprachschule</i>	'SPRACH-schu-le	Betonung auf erstem Glied signalisiert: ein Wort, keine Phrase
* <i>füreisammen</i>	<i>vereinsamen</i>	ver-' EIN-sa-men	Unbetontes Präfix {ver-}: gehört zur Verbform, nicht eigenständig

Anhand der schriftsprachlichen Muster werden auch die Strukturen der gesprochenen Sprache zugänglich und transparent gemacht, was den Lernenden ein zunehmend differenziertes Bild von der Zielsprache Deutsch ermöglicht (vgl. Pracht, 2012, S. 9).

Die longitudinale Analyse im Erwerbskorpus hat gezeigt, dass die orthographische Entwicklung nicht linear verläuft. Sie zeigt eine klare Gruppenentwicklungsdynamik, die sich jedoch deutlich von der L1-Entwicklung unterscheidet. Die Hauptschwierigkeiten verlagern sich vom ersten Semester (dominante phonografische und silbische Fehler) zum zweiten Semester hin auf komplexere syntaktische und grammatiknahe Fehler, deren Anteil signifikant zunimmt und danach weitgehend stabil bleibt. Gleichzeitig verläuft der Lernfortschritt asynchron: Kompetenzen im silbischen Bereich verbessern sich relativ früh, während Fortschritte in den komplexesten syntaktischen und grammatischen Bereichen über den Beobachtungszeitraum kaum zu verzeichnen sind. Aus dieser nicht linearen Entwicklungsdynamik lässt sich eine direkte Konsequenz für die didaktische Sequenzierung des Orthographieunterrichts ableiten. In der Anfangsphase des Erwerbs, insbesondere im ersten Semester, sollte der Unterrichtsschwerpunkt auf den fundamentalen Laut-Buchstaben-Beziehungen und den prosodischen Grundlagen des Deutschen liegen. Für fortgeschrittenere Lernende, bei denen sich die Hauptschwierigkeiten auf die syntaktische Ebene verlagern, muss der Fokus entsprechend angepasst werden. Hier sollten gezielt kontextsensitive Phänomene wie die Groß- und Kleinschreibung sowie die Getrennt- und Zusammenschreibung in den Mittelpunkt rücken. Die in dieser Studie erstellte Fehler-Rangliste bietet eine empirische Grundlage, um z. B. gezielte Rechtschreibübungen zu entwickeln, die es Lehrenden ermöglichen, den Unterricht gezielt auf die häufigsten Fehlerquellen auszurichten und somit didaktisch Schwerpunkte zu setzen.

Obwohl bereits seit 40 Jahren die curriculare Leerstelle im DaF-Unterricht bestätigt wird (vgl. Rösler, 1982), wird Orthographie im Germanistikstudium in China wie im üblichen DaF-Unterricht kaum systematisch unterrichtet. Der Lernfortschritt beruht größtenteils auf der selbstorganisierten Regelbildung der Lernenden. Während dies die beeindruckende Fähigkeit zur Selbstregulation zeigt, führt es auch zu den beobachteten, langanhaltenden Schwierigkeiten in diesen anspruchsvollen Bereichen. Es bedarf daher einer stärkeren Verankerung der Orthographie im DaF-Curriculum. Dies bedeutet jedoch kein isoliertes Regelvermitteln; vielmehr legen die Ergebnisse eine integrierte Orthographiedidaktik nahe. Diese muss orthographische Phänomene stets im Zusammenhang mit den hier diskutierten Ebenen vermitteln: der Phonologie (insb. Prosodie),

der Graphematik, der Morphologie, der Syntax sowie dem individuellen sprachlichen Repertoire der Lernenden, um so die konzeptuellen Grundlagen für eine tiefgreifende Rechtschreibkompetenz zu schaffen.

6.3 Methodische Limitationen der Studie

Die vorliegende Studie versuchte, anhand korpuslinguistischer Methoden und im Lichte der Multikompetenzhypothese sowie der CDST die Profile und Prozesse des Orthographieerwerbs chinesischer Deutschlernender zu skizzieren. Trotz größtmöglicher Sorgfalt im Forschungsdesign unterliegt die Arbeit naturgemäß einigen Limitationen, die im Folgenden diskutiert und als Anknüpfungspunkte für zukünftige Forschung verstanden werden.

Aufgrund des Studiendesigns wurde die Rechtschreibkompetenz primär anhand von Satzdiktaten erfasst. Diktate sind zwar eine etablierte und ökonomische Methode zur Erhebung orthographischer Leistungen (vgl. Blatt & Frahm, 2013), bringen jedoch spezifische Einschränkungen mit sich. So kann nur die korrekte Schreibung jener Phänomene überprüft werden, die im Diktattext auch tatsächlich vorkommen; seltenere oder komplexere Strukturen bleiben möglicherweise unberücksichtigt (vgl. Thomé & Thomé, 2020, S. 11). Zudem spielt die Aussprache der diktierenden Person eine nicht zu unterschätzende Rolle als potenzielle Störvariable. Obwohl die Aussprache im Rahmen der Analyse an einigen Stellen thematisiert (z. B. <r> im *Garten* wurde /r/ ausgesprochen) wurde, konnte eine systematische, quantitative Untersuchung der Korrelation zwischen der spezifischen prosodischen Realisierung des Inputs und der orthographischen Performanz im Rahmen des vorliegenden Forschungsdesigns nicht geleistet werden. Eine solche Analyse, die akustische Messungen des Diktat-Inputs mit den Fehlerdaten verknüpft, stellt ein wichtiges Desiderat für Anschlussuntersuchungen dar.

Eng mit dieser textuellen Datengrundlage verbunden ist eine weitere Limitation bei der Interpretation der quantitativen Ergebnisse: Aus methodischer Sicht muss einschränkend angemerkt werden, dass in der vorliegenden Studie auf die Berechnung einer exakten Basisrate (Relation der absoluten Fehlerzahlen zu den tatsächlichen Fehlerpotenzialen bzw. kritischen Punkten in den Texten) verzichtet werden musste. Aufgrund der heterogenen Zusammensetzung des DeChiLKo-Korpus, das neben den standardisierten Prüfungsaufgaben auch vielfältige Unterrichtsübungen umfasst, ließ sich die Gesamtzahl der möglichen Fehlerquellen nachträglich nicht für alle Textsorten einheitlich operationalisieren. Die dargestellten Fehlerhäufigkeiten sind daher stets vor dem Hintergrund der jeweiligen Textvorgaben zu betrachten, da nicht gänzlich ausgeschlossen werden kann, dass die Frequenz bestimmter Fehlertypen teilweise auch durch die Häufigkeit der entsprechenden Phänomene im Testmaterial bedingt ist.

Kritisch zu reflektieren ist ebenfalls die Zusammensetzung der Stichproben und das Design der beiden Subkorpora. Beide Korpora umfassen ausschließlich Germanistikstudierende. Diese Gruppe ist zwar für die Untersuchung universitärer Lernprozesse ideal, ihre Lernerfahrungen (intensiver Unterricht, hohe Motivation) sind jedoch nicht ohne Weiteres auf andere Gruppen von DaF-Lernenden (z. B. in Integrations-

kursen oder im schulischen Kontext) übertragbar. Die Generalisierbarkeit der Befunde ist somit auf diese spezifische Population beschränkt.

Das große Prüfungskorpus liefert als Querschnittsdatensatz nur eine Momentaufnahme (am Ende des zweiten Studienjahrs) des jeweiligen Entwicklungsstandes. Welche Fallstricke eine reine Querschnittsanalyse birgt, wurde im Rahmen der Diskussion des Faktors *Prüfungsstatus* deutlich: Die hier beobachtete Variabilität ist kein Beleg für eine Entwicklungsregression, sondern ein methodischer Selektionseffekt, da es sich bei den Prüfungswiederholer*innen um eine tendenziell schwächere Ausgangskompetenz handelt. Erst die Kombination mit den Längsschnittdaten des Erwerbskorpus, der echte Entwicklungsverläufe abbildet, ermöglichte eine konkrete Interpretation. Die geringe Probandenzahl ($n = 5$) im Erwerbskorpus schränkt dessen statistische Aussagekraft jedoch naturgemäß ein. Zukünftige Forschung sollte idealerweise auf größere Längsschnittkorpora zurückgreifen.

Eine weitere wesentliche Einschränkung liegt in der unzureichenden Erfassung von soziodemografischen und kognitiven Faktoren. In der vorliegenden Studie wurden diese individuellen Lernervariablen nicht direkt erhoben, sondern nur indirekt über die groben Faktoren *Hochschulkategorie* und *Region* approximiert. Diese Kategorien können jedoch die erhebliche Heterogenität innerhalb der Lernergruppen nicht abbilden. Individuelle kognitive Voraussetzungen wie das phonologische Bewusstsein, das Kurzzeitgedächtnis, die Abrufgeschwindigkeit und die visuelle Aufmerksamkeit (vgl. Mannhaupt, 2003) haben einen nachgewiesenen signifikanten Einfluss auf den Orthographieerwerb (vgl. u. a. Bulut, 2018). Da diese Daten nicht erfasst wurden, konnte ihr spezifischer Beitrag zum orthographischen Lernerfolg nicht konkret analysiert werden. Ein direkter Vergleich der Befunde mit anderen Studien ist somit unmöglich. Ein erheblicher Teil der Varianz in den Lernleistungen bleibt somit unerklärt. Zukünftige Studien sollten daher gezielte Erhebungsinstrumente wie standardisierte Tests oder detaillierte Fragebögen einsetzen, um ein präziseres Bild der lernenden Person zu erhalten und die Wechselwirkungen zwischen individuellen Anlagen und dem Orthographieerwerb genauer zu untersuchen.

Eine weitere Limitation liegt in der Natur der Fehleranalyse selbst. Die Unterscheidung zwischen Kompetenz- und Performanzfehlern (vgl. Corder, 1967, 161 ff.) ist anhand der Fehleranalyse von Korpusdaten nur schwer zu operationalisieren. Obwohl die Stabilität der Leistung über verschiedene Aufgabentypen hinweg auf eine relativ gefestigte orthographische Kompetenz hindeutet, lässt sich im Einzelfall nie gänzlich ausschließen, dass ein Fehler auf mangelnde Konzentration zurückzuführen ist.

Schließlich ist als Limitation anzuführen, dass im Forschungsdesign keine Kontroll- bzw. Vergleichsgruppen vorgesehen waren. Die Hypothesen zum spezifischen Einfluss der L1 Chinesisch wurden durch kontrastive Analysen und den Vergleich mit bestehender Forschung zu anderen Lernergruppen gestützt. Ein direkter Vergleich der orthographischen Leistungen von chinesischen Lernenden mit denen von DaF-Lernenden mit anderen L1 (z. B. einer romanischen oder slawischen Sprache) oder mit deutschen Muttersprachler*innen innerhalb desselben Untersuchungsdesigns würde jedoch eine noch validere Überprüfung der L1-spezifischen Hypothesen ermöglichen.

7 Schluss

Abschließend lässt sich festhalten, dass die vorliegende Studie den orthographischen Erwerb chinesischer Deutschlernender als ein komplexes, dynamisches System beleuchtet hat. Die Ergebnisse deuten darauf hin, dass die spezifischen orthographischen Herausforderungen der Lernenden aus einem vielschichtigen Zusammenspiel von L1-Einflüssen (Chinesisch/*Pinyin*), L2-Transfereffekten (Englisch) und den inhärenten Schwierigkeiten des deutschen Schriftsystems resultieren. Die Entwicklung der Rechtschreibkompetenz verläuft dabei nicht linear, sondern ist durch bereichsspezifische, gestaffelte Fortschritte sowie eine hohe individuelle Variabilität gekennzeichnet. Beeinflusst wird dieser Prozess maßgeblich durch die Lerndauer sowie durch externe, institutionelle Rahmenbedingungen wie die Hochschulkategorie und die Region, die als Indikatoren für die Lernvoraussetzungen und die Qualität der Lernumgebung fungieren.

Die methodische Grundlage dieser Arbeit – die Erstellung des Deutsch-Chinesischen Lernerkorpus (DeChiLKo) mit seinem detaillierten, mehrschichtigen Annotationsschema – ermöglichte es, über eine reine Fehlerzählung hinauszugehen und die qualitativen Muster systematisch zu analysieren. Die detaillierte Fehleranalyse hat die zentrale Rolle des phonographischen Bereichs zwar quantitativ bestätigt, zugleich aber offenbart, dass die tieferliegenden Ursachen oft konzeptueller Natur sind. So sind die Schwierigkeiten bei der Getrennt- und Zusammenschreibung keine isolierten Regelverstöße, sondern offenbaren vielmehr eine grundlegende Herausforderung: die korrekte syntaktische Segmentierung. Diese Schwierigkeiten lassen sich auch auf die unterschiedlichen prosodischen Betonungsmuster zurückführen. Die vorliegende Arbeit konnte somit nachweisen, dass eine rein phonographisch orientierte Fehleranalyse zu kurz greift und erst die Betrachtung des gesamten orthographischen Systems die tatsächlichen Lernhürden sichtbar macht.

Diese Befunde stützen den gewählten integrierten Bezugsrahmen aus Multi-Kompetenz-Hypothese und der Theorie komplexer dynamischer Systeme (CDST). Die beobachteten Transfereffekte aus dem *Pinyin*-System und der L2 Englisch bestätigen, dass die Lernenden beim Orthographieerwerb auf ihr gesamtes multikompetentes Repertoire zurückgreifen. Die Entwicklungsdynamik erweist sich als nicht linear und ist von individueller Variabilität sowie gestaffelten Lernfortschritten geprägt. Sie unterscheidet sich damit signifikant von Erwerbsmodellen für Lernende mit Deutsch als L1 und veranschaulicht eindrucksvoll die Prinzipien eines sich selbst organisierenden, dynamischen Systems. Dieser Prozess wird, wie die statistische Analyse darlegte, maßgeblich von institutionellen und regionalen Rahmenbedingungen moderiert, die die Lernumgebung und die individuellen Entwicklungspfade prägen.

Für die Praxis folgt daraus die Notwendigkeit einer prozessorientierten, diagnostisch fundierten und mehrsprachigen Orthographiedidaktik. Die orthographischen Abweichungen chinesischer Lernender sind als emergente Phänomene zu verstehen, die einer differenzierten qualitativen Analyse auf linguistischer Grundlage bedürfen. An-

statt einer reinen Fehlerkorrektur erfordert dies eine Vermittlung, die insbesondere die für Lernende mit einer Tonsprache schwer zugänglichen prosodischen Grundlagen der deutschen Orthographie explizit thematisiert.

Die vorliegende Arbeit hat die Komplexität des L3-Orthographieerwerbs aufgezeigt und ein systemisches Modell zur Erklärung der Befunde angeboten. Gleichzeitig eröffnen die Ergebnisse und die methodischen Limitationen neue Perspektiven für weiterführende Forschung. Ein erster wichtiger Schritt wäre die Erweiterung des korpuslinguistischen Untersuchungsdesigns, um die hier gewonnenen Erkenntnisse zu validieren und zu generalisieren. Um die Rechtschreibleistung mehrsprachiger Lernenden noch besser einordnen zu können, wären vergleichende Studien wünschenswert, die explizit Kontrollgruppen (z. B. deutsche Muttersprachler*innen oder DaF-Lerner*innen mit anderen Erstsprachen) einbeziehen. Darüber hinaus könnte die systematische Erfassung zusätzlicher Variablen wie spezifischer Lehrmethoden an den Universitäten oder der Vergleich unterschiedlicher Textsorten (z. B. Diktate vs. freie Aufsätze) helfen, das relative Gewicht der verschiedenen Einflussfaktoren auf die orthographische Leistung genauer zu bestimmen.

Ein zweiter, fundamentaler Forschungsausblick betrifft die Analyse der kognitiven Prozesse, die der orthographischen Performanz zugrunde liegen. Während die vorliegende Korpusstudie das *Ergebnis* eines Schreibprozesses (den Diktattext) analysierte, wäre es für ein tieferes Verständnis der Lernermechanismen unerlässlich, den *Prozess selbst* zu untersuchen. Hierfür bieten sich prozessorientierte Erhebungsmethoden an. Mittels „Lautem Denken“ oder Rechtschreibgesprächen (Geist & Hesselbarth, 2017) könnten die orthographischen Entscheidungs- und Problemlösungsstrategien der Lernenden direkt erfasst werden. Ergänzend könnte die Analyse der Handschriftndynamik mittels digitaler Stifte und Tablets Aufschluss über die kognitive Anstrengung während des Schreibens geben (vgl. Angelillo et al., 2019). Diese Methode erfasst Parameter wie die Dauer und Position von Pausen, die als verlässliche Indikatoren für die kognitive Belastung dienen (vgl. Hurschler Lichtsteiner, Wicki & Falmann, 2018). Beispielsweise können längere und häufigere Pausen, insbesondere innerhalb von Wörtern, auf eine höhere kognitive Belastung hindeuten, die oft mit Rechtschreibschwierigkeiten zusammenhängt (vgl. Mazur & Quignard, 2025). Solche prozessualen Daten würden es ermöglichen, die in dieser Arbeit postulierten Dynamiken, wie den Rückgriff auf Attraktorzustände unter kognitiver Last, direkt in der Echtzeit-Performanz zu beobachten.

Die Kombination aus erweiterten korpuslinguistischen Designs und prozessorientierten Methoden verspricht, unser Verständnis des komplexen dynamischen Systems des L3-Orthographieerwerbs weiter zu vertiefen und die hier entwickelten Modelle zu überprüfen und zu verfeinern.

Literaturverzeichnis

- Angelillo, M. T., Impedovo, D., Pirlo, G., Sarcinella, L. & Vessio, G. (2019). Handwriting Dynamics as an Indicator of Cognitive Reserve: An Exploratory Study. In *2019 IEEE International Conference on Systems, Man and Cybernetics (SMC)* (S. 835–840). IEEE. <https://doi.org/10.1109/SMC.2019.8914157>
- Anleitungskomitee für den Fremdsprachenunterricht an Hochschulen des chinesischen Bildungsministeriums. (2013). *Prüfungsordnungen für das Germanistik-Grundstudium und -Hauptstudium im Hochschulwesen Chinas*. [Gāoděng xuéxiào déyǔ zhuānyè sì, bā jí kǎoshì dàgāng]. Shanghai: Verlag für Fremdsprachenausbildung in Shanghai.
- Aronin, L. & Hufeisen, B. (Hrsg.). (2009). *The Exploration of Multilingualism. Development of Research on L3, Multilingualism and Multiple Language Acquisition*. Amsterdam: John Benjamins. <https://doi.org/10.1075/aals.6>
- Auswärtiges Amt, Goethe-Institut, Deutsche Welle, DAAD & ZfA (Hrsg.) (2020): Deutsch als Fremdsprache weltweit. Datenerhebung 2020. Berlin/München: Auswärtiges Amt.
- Bai, C., Chi, W. & Qian, X. (2014). Do college entrance examination scores predict undergraduate GPAs? A tale of two universities. *China Economic Review*, 30, 632–647. <https://doi.org/10.1016/j.chieco.2013.08.005>
- Bailey, N., Madden, C. & Krashen, S. D. (1974). Is There a „Natural Sequence“ in Adult Second Language Learning? *Language Learning*, 24(2), 235–243. <https://doi.org/10.1111/j.1467-1770.1974.tb00505.x>
- Bates, D., Mächler, M., Bolker, B. & Walker, S. (2015). Fitting Linear Mixed-Effects Models Using lme4. *Journal of Statistical Software*, 67(1), 1–48. <https://doi.org/10.18637/jss.v067.i01>
- Bausch, K.-R. & Kasper, G. (1979). Der Zweitsprachenerwerb: Möglichkeiten und Grenzen der „großen“ Hypothesen. *Linguistische Berichte*, 64, 3–35.
- Bassetti, B. (2017). Orthographic input and phonological acquisition. In L. Slabakova (Ed.), *Inquiries in linguistic development*. John Benjamins.
- Becker, T. (2013). *Schriftspracherwerb in der Zweitsprache*. Zugl.: Dortmund, Techn. Univ., Habil.-Schr., 2011. Schneider-Verl. Hohengehren, Baltmannsweiler.
- Becker, T. & Siekmann, K. (2012). Diagnose orthographischer Fähigkeiten bei mehrsprachigen Kindern. In W. Griefhaber & Z. Kalkavan-Aydın (Hrsg.), *Orthographie- und Schriftspracherwerb bei mehrsprachigen Kindern* (1. Aufl., S. 169–188). Stuttgart: Fillibach Verlag; Klett.
- Beeh, C., Drewnowska-Vargáné, E., Kappel, P., Modrián-Horváth, B., Nolda, A., Rauzs, O. et al. (2021). *Dulko-Handbuch. Aufbau und Annotationsverfahren des deutsch-ungarischen Lernerkorpus. Version 1.0*. Szeged: Institut für Germanistik der Universität Szeged. <https://doi.org/10.14232/dulko-handbuch-v1.0>
- Benedix, A. (2009). *Chinesisch als Fremdsprache in der Sekundarstufe. Binnendifferenzierung und die Gestaltung von Unterrichtsmaterialien*. Zugl.: Bochum, Univ., Diss., 2007. Marburg: Tectum-Verl.

- Berg, K., Hartmann, S. & Claeser, D. (2023). Are some morphological units more prone to spelling variation than others? A case study using spontaneous handwritten data. *Morphology*, (34), 173–188. <https://doi.org/10.1007/s11525-023-09417-4>
- Bildungsministerium. (2006). 高等学校德语专业德语本科教学大纲 *Rahmenplan für das Bachelorstudium im Fach Deutsch an Hochschulen und Universitäten in China*. Shanghai: Verlag für Fremdsprachenausbildung in Shanghai.
- Birkel, P. (2009). Rechtschreibleistung im Diktat – eine objektiv beurteilbare Leistung? *Didaktik Deutsch*, 27, 5–32. <https://doi.org/10.25656/01:21337>
- Blatt, I. (2010). Sprachsystematische Rechtschreibdidaktik: Konzept, Materialien, Tests. In G. Hinney (Hrsg.), *Schriftsystem und Schrifterwerb: linguistisch, didaktisch, empirisch* (Reihe germanistische Linguistik, v.289, 1. Aufl., S. 101–132). Berlin: De Gruyter.
- Blatt, I. & Frahm, S. (2013). Explorative Analysen zur Entwicklung der Rechtschreibkompetenz im Rahmen der NEPS-Studie (Klassenstufe 5–7). *Didaktik Deutsch: Halbjahresschrift für die Didaktik der deutschen Sprache und Literatur*, 18, 13–36. <https://doi.org/10.25656/01:21180>
- Blatt, I., Hein, C. & Streubel, B. (2015). *Entdecke die Schrift. Schreiben und Lesen lernen mit dem Bärenboot*. München: Oldenbourg Schulbuchverlag.
- Bredel, U. (2006). Die Herausbildung des syntaktischen Prinzips in der Historiogenese und in der Ontogenese der Schrift. In U. Bredel, H. Günther, P. Klotz, J. Ossner & G. Siebert-Ott (Hrsg.), *Didaktik der deutschen Sprache* (S. 139–163). Paderborn: Schöningh. <https://doi.org/10.1515/9783110921199.139>
- Bredel, U. (2013). Silben und Füße im Deutschen und Türkischen – Orthographieerwerb des Deutschen durch türkischsprachige Lerner/innen. In A. Deppermann (Hrsg.), *Das Deutsch der Migranten* (S. 369–390). Berlin: De Gruyter. <https://doi.org/10.1515/9783110307894.369>
- Bredel, U., Fuhrhop, N. & Noack, C. (2011). *Wie Kinder lesen und schreiben lernen*. Tübingen: Francke Verlag.
- Bredel, U. & Müller, A. (2010). *Schriftsystem und Schrifterwerb: linguistisch, didaktisch, empirisch* (Reihe germanistische Linguistik, v.289, 1. Aufl.). Berlin: De Gruyter. <https://doi.org/10.1515/97831102322571>
- Bredel, U. & Röber, C. (2011). Zur Gegenwart des Orthographieunterrichts. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5, S. 3–9). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Brinkmann, E. (2003). „Farrat da war nichz Schwirich ...“. In E. Brinkmann, N. Kruse & C. Osburg (Hrsg.), *Kinder schreiben und lesen. Beobachten-verstehen-lehren* (S. 147–154). Freiburg: Fillibach Verlag. <https://doi.org/10.25656/01:16953>
- Brügelmann, H. & Brinkmann, E. (1994). Stufen des Schriftsprachenerwerbs und Ansätze zu seiner Förderung. In H. Brügelmann & S. Riegler (Hrsg.), *Wie wir recht schreiben lernen* (S. 44–52). Lengwil: Libelle.
- Budde, M., Riegler, S. & Wiprächtiger-Geppert, M. (2011). *Sprachdidaktik* (Akademie-Studienbücher). Berlin: Akademie Verlag. <https://doi.org/10.1524/9783050052823>

- Bülow, L., de Bot, K. & Hilton, N. (2017). Zum Nutzen der Complex Dynamic Systems Theory (CDST) für die Erforschung von Sprachvariation und Sprachwandel. In H. Christen, P. Gilles & C. Purschke (Hrsg.), *Räume, Grenzen, Übergänge: Akten des 5. Kongresses der Internationalen Gesellschaft für* (S. 45–69). Stuttgart: Franz Steiner Verlag.
- Bulut, N. (2018). *Individuelle Rechtschreibentwicklung. Eine Längsschnittuntersuchung zur Bedeutung von Einflussfaktoren auf die Wortschreibung* (Thema Sprache – Wissenschaft für den Unterricht Band 27). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Cameron, A. & Trivedi, P. (2013). *Regression Analysis of Count Data*. Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9781139013567>
- Cheung, H., Chen, H. C., Lai, C. Y., Wong, O. C. & Hills, M. (2001). The development of phonological awareness: effects of spoken language experience and orthography. *Cognition*, 81(3), 227–241. [https://doi.org/10.1016/S0010-0277\(01\)00136-6](https://doi.org/10.1016/S0010-0277(01)00136-6)
- Chi, X. (2016). *Die Fehleranalyse im DaF-Unterricht. Untersuchung von Aufsätzen chinesischer Germanistikstudenten*. Bachelorarbeit.
- Chiao, W. J. & Kelz, H. P. (1985). *Chinesische Aussprache. Ein Lernprogramm = Hanyu fayin* (Dümmlerbuch, Bd. 6101, 2. Aufl.). Bonn: Dümmler.
- Chomsky, N. (1965). *Aspects of the theory of syntax*. Cambridge: M. I. T. Press. <https://doi.org/10.21236/AD0616323>
- Chomsky, N. (1985). *Knowledge of Language: Its Nature, Origin, and Use*. New York: Praeger Publishers. Verfügbar unter: http://www.thatmarcusfamily.org/philosophy/Course_Websites/Readings/Chomsky%20-%20Knowledge%20of%20Language.pdf (13.06.2026)
- Clahsen, H. (1982). *Spracherwerb in der Kindheit. Eine Untersuchung zur Entwicklung der Syntax bei Kleinkindern*. Tübingen: Narr Francke Attempto.
- Cook, V. (1992). Evidence for multi-competence. *Language Learning*, 42(4), 557–591. <https://doi.org/10.1111/j.1467-1770.1992.tb01044.x>
- Cook, V. (2012). 'Multi-competence'. In C. A. Chapelle (Hrsg.), *The Encyclopedia of Applied Linguistics* (S. 3768–3774). New York: Wiley-Blackwell.
- Cook, V. (2016). Premises of multi-competence. In V. Cook & W. Li (Eds.), *The Cambridge Handbook of Linguistic Multi-Competence* (Cambridge Handbooks in Language and Linguistics, 1st ed., S. 1–25). Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9781107425965.001>
- Corder, S. P. (1967). The Significance of Learner's Errors. *International Review of Applied Linguistics in Language Teaching*, 5, 161–170. <https://doi.org/10.1515/iral.1967.5.1-4.161>
- Corder, S. P. (1986). *Error analysis and interlanguage* (4. impr.). Oxford: Oxford University Press.
- Corvacho Del Toro, I. M. (2004). *Zweitalphabetisierung und Orthographieverwerb. Deutsch-spanisch bilinguale Kinder auf dem Weg zur biliteralen Kompetenz* (Europäische Hochschulschriften: Reihe 1, Deutsche Sprache und Literatur, Bd. 1885). Frankfurt am Main, Berlin, Bern, Bruxelles, New York, Oxford, Wien: Lang Peter GmbH, Internationaler Verlag der Wissenschaften.

- Corvacho Del Toro, I. M. (2013). *Fachwissen von Grundschullehrkräften: Effekt auf die Rechtschreibleistung von Grundschulern*. Bamberg: University of Bamberg Press.
<https://doi.org/10.20378/irb-1538>
- Corvacho Del Toro, I. M. (2016). Zur qualitativen Rechtschreibfehleranalyse und einer schriftsystematischen lernförderlichen Behandlung der Rechtschreibstörung. *Zeitschrift für Kinder- und Jugendpsychiatrie und Psychotherapie* [A qualitative analysis of spelling mistakes and a systematic supportive learning instruction of spelling disorder], 44(5), 397–408. <https://doi.org/10.1024/1422-4917/a000457>
- Corvacho Del Toro, I. M. & Lorenz, J. H. (2021). Fehler: Symptom und Lernchance zugleich. *Lernen und Lernstörungen*, 10(1), 1–2. <https://doi.org/10.1024/2235-0977/a000321>
- Cremerius, R. (Autor). (2012). *Aussprache und Schrift des Chinesischen. Eine Einführung*. Hamburg: Buske.
- De Bot, K. (2017). Complexity Theory and Dynamic Systems Theory: Same or different? In *Complexity Theory and Language Development*. In *Celebration of Diane Larsen-Freeman*. (S. 11–58). Amsterdam: John Benjamins. <https://doi.org/10.1075/lllt.48.03deb>
- De Bot, K. & Larsen-Freeman, D. (2011). Researching Second Language Development from a Dynamic Systems Theory perspective. In M. Verspoor, K. de Bot & W. Lowie (Eds.), *A dynamic approach to second language development methods and techniques* (Language learning & language teaching (LL<), vol. 29, S. 5–24). Amsterdam: John Benjamins. <https://doi.org/10.1075/lllt.29.01deb>
- De Bot, K., Lowie, W. & Verspoor, M. (2007). A dynamic systems theory approach to second language acquisition. *Bilingualism: Language and Cognition*, 10(1), 7–21. <https://doi.org/10.1017/S1366728906002732>
- De Bot, K. & Makoni, S. (2005). *Language and Aging in Multilingual Contexts*. Tallinn: Multilingual Matters. <https://doi.org/10.21832/9781853598425>
- DESI-Konsortium. (2008). *Unterricht und Kompetenzerwerb in Deutsch und Englisch. Ergebnisse der DESI-Studie*. Weinheim: Beltz Pädagogik. <https://doi.org/10.25656/01:3149>
- Deutscher Akademischer Austauschdienst. (2019). *China: Daten & Analysen zum Hochschul- und Wissenschaftsstandort. (DAAD-Bildungssystemanalyse)*. Bonn: Deutscher Akademischer Austauschdienst. Verfügbar unter: https://www.chinazentren.de/wp-content/uploads/2020/02/DAAD_Bildungssystemanalyse2019.pdf (13.06.2026)
- Diehl, E. (2000). *Grammatikunterricht: Alles für die Katz? Untersuchungen zum Zweitspracherwerb Deutsch* (Reihe germanistische Linguistik, Bd. 220). Tübingen: Max Niemeyer Verlag.
- Ding, L. (2018). *Analyse der lexikalischen Fehler in Aufsätzen der PGG – eine korpusbasierte Untersuchung*. Magisterarbeit. Schanghai International Studies University, Shanghai.
- Drosdowski, G. & Eisenberg, P. (1995). *Duden Grammatik der deutschen Gegenwartssprache* (Der Duden in 12 Bänden, Bd. 4, 5. völlig neu bearb. und erw. Aufl.). Mannheim: Dudenverlag.
- Duden (o. J.). Die Verteilung der Wortarten im Rechtschreibduden. <https://www.duden.de/sprachwissen/sprachratgeber/Die-Verteilung-der-Wortarten-im-Rechtschreibduden> (12.06.2026)

- Dulay, H. & Burt, M. (1973). Should we teach Children Syntax? *Language Learning*, 23(3), 245–258. <https://doi.org/10.1111/j.1467-1770.1974.tb00234.x>
- Dulay, H. & Burt, M. (1974). Natural sequences in child second language acquisition. *Language Learning*, 24(1), 37–53. <https://doi.org/10.1111/j.1467-1770.1974.tb00234.x>
- DWDS Kernkorpus 21. *Textkorpus bereitgestellt durch das Digitale Wörterbuch der deutschen Sprache*. Verfügbar unter: <https://www.dwds.de/d/korpora/kern21> (13.06.2026)
- Eberhard, David M., Gary F. Simons & Alison J. Robinson (Hrsg.). 2026. *Ethnologue: Languages of the World*. Twenty-ninth edition. Dallas, Texas: SIL Global. <https://www.ethnologue.com/insights/most-spoken-language> (12.06.2026)
- Edmondson, W. J. & House, J. (2011). *Einführung in die Sprachlehrforschung* (4. überarb. Aufl.). A. Francke.
- Eichler, W. & Thomé, G. (1995). Bericht aus dem DFG-Forschungsprojekt „Innere Regelformbildung im Orthographieerwerb im Schulalter“. In H. Brügelmann & H. Balhorn (Hrsg.), *Am Rande der Schrift* (S. 35–42). Lengwil: Libelle.
- Eisenberg, P. (1995). Die deutsche Orthographie und Deutsch als Fremdsprache: analoge Strukturierung von System und Erwerb. *Jahrbuch Deutsch als Fremdsprache*, 21, 171–184. Verfügbar unter: https://www.uni-potsdam.de/fileadmin/projects/peter-eisenberg-schriften/PDFs/Eisenberg_1995_Die-deutsche-Orthographie-und-Deutsch-als-Fremdsprache.pdf (13.06.2026)
- Eisenberg, P. (2011). Grundlage der deutschen Wortschreibung. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5, S. 83–95). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Eisenberg, P. (2013). *Das Wort* (Grundriss der deutschen Grammatik, Bd. 1, 4. aktualisierte und überarb. Aufl.). Stuttgart: Metzler. <https://doi.org/10.1007/978-3-476-00757-5>
- Eisenberg, P. (2016). Phonem und Graphem. In Dudenredaktion (Hrsg.), *Duden – die Grammatik. Unentbehrlich für richtiges Deutsch* (Bd. 4). Mannheim: Dudenverlag.
- Eisenberg, P. (2017). *Deutsche Orthografie. Regelwerk und Kommentar*. Berlin, Boston: De Gruyter. <https://doi.org/10.1515/9783110525229>
- Ellis, N. C. & Larsen-Freeman, D. (2006). Language Emergence: Implications for Applied Linguistics – Introduction to the Special Issue. *Applied Linguistics*, 27(4), 558–589. <https://doi.org/10.1093/applin/aml028>
- Ellis, R. (1997). *The study of the second language acquisition* (Oxford applied linguistics, 5th impr.). Oxford: Oxford University Press.
- Fan, L. (2017). *Deutsch nach Englisch in China. Subjektive Vorstellungen Lehrender über das Deutsch-als-L3-Lehren* (Tertiärsprachen). Tübingen: Stauffenburg Verlag.
- Farley, A. & Yang, H. H. (2020). Comparison of Chinese Gaokao and Western university undergraduate admission criteria: Australian ATAR as an example. *Higher Education Research & Development*, 39(3), 470–484. <https://doi.org/10.1080/07294360.2019.1684879>
- Fay, J. (2010). *Die Entwicklung der Rechtschreibkompetenz beim Textschreiben*. Frankfurt am Main: Peter Lang. <https://doi.org/10.3726/978-3-653-00718-3>

- Feilke, H. (2011). Der Erwerb der das/dass-Schreibung. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5, S. 340–354). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Feng, C. X. (2021). A comparison of zero-inflated and hurdle models for modeling zero-inflated count data. *Journal of Statistical Distributions and Applications*, 8(8), 1–19. <https://doi.org/10.1186/s40488-021-00121-4>
- Firth, A. & Wagner, J. (1997). On Discourse, Communication, and (Some) Fundamental Concepts in SLA Research. *The Modern Language Journal*, 81(3), 285–300. <https://doi.org/10.1111/j.1540-4781.1997.tb05480.x>
- Fix, M. (2002). „Die Recht Schreibung ferbessern“. Zur Orthografischen Kompetenz in der Zweitsprache Deutsch. *Didaktik Deutsch*, 7(12), 39–55.
- Flege, J. E. (1995). Second Language Speech Learning: Theory, Findings, and Problems. In W. Strange (Hrsg.), *Speech Perception and Linguistic Experience: Issues in Cross-Language Research* (S. 233–277). New York, Timonium, MD. Zugriff am 04.11.2024. Verfügbar unter: https://www.researchgate.net/profile/James-Flege/publication/333815781_Second_language_speech_learning_Theory_findings_and_problems/links/5d071d2692851c900442d6b2/Second-language-speech-learning-Theory-findings-and-problems.pdf (13.06.2026)
- Fleischer, W. & Barz, I. (2012). *Wortbildung der deutschen Gegenwartssprache* (4. Aufl.). Berlin: De Gruyter. <https://doi.org/10.1515/9783110256659>
- Földes, C. (2000). Rechtschreibunterricht in der Lernsprache Deutsch nach der Orthografiereform. *Zielsprache Deutsch*, (31), 15–30.
- Fox-Boyer, A., Groos, I. & Schauß-Golecki, K. (2015). *Kindliche Aussprachestörungen. Ein Ratgeber für Eltern, Erzieher, Therapeuten und Ärzte* (Ratgeber für Angehörige, Betroffene und Fachleute, 3. überarb. Aufl.). Idstein: Schulz-Kirchner.
- Fries, C. C. (1945). *Teaching and Learning English as a Foreign Language*. Ann Arbor: University of Michigan Press.
- Frith, U. (1985). Beneath the surface of developmental dyslexia. In K. E. Patterson, J. C. Marshall & M. Coltheart (Hrsg.), *Surface Dyslexia* (S. 301–329). Oxford: Routledge Taylor & Francis Group. <https://doi.org/10.4324/9781315108346-18>
- Fuhrhop, N. (2000). Zeigen Fugenelemente die Morphologisierung von Komposita an? In R. Thieroff, M. Tamrat, N. Fuhrhop & O. Teuber (Hrsg.), *Deutsche Grammatik in Theorie und Praxis* (S. 201–214). Berlin: De Gruyter. <https://doi.org/10.1515/9783110933932.201>
- Fuhrhop, N. (2007). *Zwischen Wort und Syntagma. Zur grammatischen Fundierung der Getrennt- und Zusammenschreibung*. Tübingen: Max Niemeyer Verlag. <https://doi.org/10.1515/9783110936544>
- Fuhrhop, N. (2010). Getrennt- und Zusammenschreibung: Kern und Peripherie. Recht-schreibdidaktische Konsequenzen aus dieser Unterscheidung. In G. Hinney (Hrsg.), *Schriftsystem und Schrifterwerb: linguistisch, didaktisch, empirisch* (Reihe germanistische Linguistik, v.289, 1. Aufl., S. 235–258). Berlin: De Gruyter.

- Fuhrhop, N. (2011). System der Getrennt- und Zusammenschreibung. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5, S. 107–128). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Fuhrhop, N. (2021). *Orthografie* (5. Aufl.). Heidelberg: Universitätsverlag Winter. <https://doi.org/10.33675/978-3-8253-7286-6>
- Fuhrhop, N. (2022). Vergleichende Graphematik als Ressource für den mehrsprachigen Schriftspracherwerb – Zur unterschiedlichen Nutzung des weitgehend gleichen Formeninventars. In K. Nimz, C. Noack & K. Schmidt (Hrsg.), *Mehrsprachigkeit und Orthographie. Empirische Studien an der Schnittstelle von Linguistik und Sprachdidaktik* (Thema Sprache – Wissenschaft für den Unterricht, S. 15–50). Bielefeld: wbv Publikation.
- Fuhrhop, N. & Peters, J. (2013). *Einführung in die Phonologie und Schrift*. Stuttgart: Metzler. <https://doi.org/10.1007/978-3-476-00597-7>
- Fuhrhop, N. & Romstadt, J. (2021). Orthographiefehler im Abitur – Eine sprachwissenschaftliche Bestandsaufnahme. In M. Kepser, S. Schallenberger & H.-G. Müller (Hrsg.), *Neue Wege des Orthografieerwerbs. Forschung – Vermittlung – Reflexion* (1. Auflage, S. 189–208). Wien: Lemberger Publishing.
- Gabrič, A. (2019). Analyse der Fehler in Texten der Deutschlernenden auf der Anfängerstufe. *Journal for Foreign Languages*, 11(1), 173–190. <https://doi.org/10.4312/vestnik.11.173-190>
- Geist, B. & Hesselbarth, C. (2017). „also <h> ist groß, weil das substantiv ist?“ Rechtschreibfähigkeit und Rechtschreibgespräche einer Schülerin mit Deutsch als Zweitsprache. In R. Sigel & E. Inckemann (Hrsg.), *Diagnose und Förderung von Kindern mit Zuwanderungshintergrund im Sprach- und Schriftspracherwerb. Theorien, Konzeptionen und Methoden in den Jahrgangstufen 1 und 2 der Grundschule* (S. 137–148). Bad Heilbrunn: Verlag Julius Klinkhardt.
- Geschäftsstelle des Rats für deutsche Rechtschreibung. (2024). *Amtlisches Regelwerk der deutschen Rechtschreibung. Regeln und Wörterverzeichnis*. Mannheim: IDS-Verlag.
- Glück, H. & Rödel, M. (Hrsg.). (2016). *Metzler Lexikon Sprache* (5. aktualisierte und überarbeitete Auflage). Stuttgart: J. B. Metzler Verlag. <https://doi.org/10.1007/978-3-476-05486-9>
- Goldberger, J. (2017). *Chinas Hochschulen im Weltbildungssystem: Analyse von Internationalisierungsstrategien*. Diss. Humboldt-Universität zu Berlin, Berlin.
- Gombert, J. E. & Fayol, M. (1992). Writing in preliterate children. *Learning and Instruction*, 2(1), 23–41. [https://doi.org/10.1016/0959-4752\(92\)90003-5](https://doi.org/10.1016/0959-4752(92)90003-5)
- Granger, S., Gilquin, G. & Meunier, F. (2015). Introduction: learner corpus research – past, present and future. In S. Granger, G. Gilquin & F. Meunier (Hrsg.), *The Cambridge Handbook of Learner Corpus Research* (S. 1–9). Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9781139649414.001>
- Granger, S. & Paquot, M. (2017). *Core metadata for learner corpora. Draft 1.0*. Unveröffentlichtes Manuskript, Louvain-la-Neuve.
- Gregg, K. R. (2010). Shallow draughts: Larsen-Freeman and Cameron on complexity. *Second Language Research*, 26(4), 549–560.

- Grießhaber, W. (2004). Einblicke in zweitsprachliche Schrifterwerbsprozesse. In M. Baumann & J. Ossner (Hrsg.), *Diagnose und Schrift II: Schreibfähigkeiten*. (Osnabrücker Beiträge zur Sprachtheorie, 67/04, S. 65–87). Osnabrück: Osnabrücker Beiträge zur Sprachtheorie. <https://doi.org/10.1017/CBO9781139649414.001>
- Grießhaber, W. (2016). *Spracherwerbsprozesse in Erst- & Zweitsprache. Eine Einführung* (3. unveränderte Auflage). Duisburg: Universitätsverlag Rhein-Ruhr.
- Grießhaber, W. & Kalkavan-Aydın, Z. (Hrsg.). (2012). *Orthographie- und Schriftspracherwerb bei mehrsprachigen Kindern* (1. Aufl.). Stuttgart: Fillibach Verlag; Klett.
- Günther, H. [Hartmut]. (2010). *Beiträge zur Didaktik der Schriftlichkeit* (Kölner Beiträge zur Sprachdidaktik Reihe A, Bd. 6). Duisburg: Gilles & Francke. Verfügbar unter: <http://www.koebes.uni-koeln.de/koebes06.pdf> (13.06.2026)
- Günther, H. [Hartmut] & Gaebert, D.-K. (2011). Das System der Groß- und Kleinschreibung. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Günther, K. B. (1986). Ein Stufenmodell kindlicher Lese- und Schreibstrategien. In H. Brügelmann (Hrsg.), *ABC und Schriftsprache* (S. 32–54). Konstanz: Faude.
- Hachenberg, P. (2003). Probleme chinesischer Lerner mit dem deutschen Vokalsystem. In U. Wannagat (Hrsg.), *Deutsch als zweite Fremdsprache in Ostasien* (Arbeiten zur angewandten Linguistik, Bd. 2, S. 97–128). Tübingen: Stauffenburg Verlag.
- Halman, S. (2018). *Alphabetisierungskurs Sekundarstufe I. Materialien für den Schriftspracherwerb – Wortebene* (Sekundarstufe I, 1. Auflage). Augsburg: Auer.
- Hanisch, A. (2014). Das Rechtschreibinterview als Diagnoseinstrument. Mit Kindern über ihr Rechtschreibwissen sprechen. *Grundschule Deutsch*, 44, 6–8.
- Hans-Bianchi, B. (2007). Der italienische Weg zur deutschen Rechtschreibung. Überlegungen zu Schreibprozess und Schreiberwerb in der Fremdsprache Deutsch. *DaF-Werkstatt*, (9/10), 77–96.
- Hans-Bianchi, B. & Katelhön, P. (2010). Orthographieerwerb in der L3. *Zeitschrift für Fremdsprachenforschung*, 21(1), 31–51.
- Hayes-Harb, R., Brown, K. & Smith, B. L. (2018). Orthographic Input and the Acquisition of German Final Devoicing by Native Speakers of English. *Language and Speech*, 61(4), 547–564. <https://doi.org/10.1177/0023830917710048>
- Heller, K. (IDS Mannheim, Hrsg.). (1996). *SPRACHREPORT Extraausgabe: Rechtschreibreform. Eine Zusammenfassung*. Verfügbar unter: <https://pub.ids-mannheim.de/lau fend/sprachreport/sr96extra.html?template=/allgemein/template/print.tpl#KH> (13.06.2026)
- Henderson, M. M. T. (1985). The Notion of Interlanguage. *Journal of Modern Language Learning*, (21), 23–27.
- Henschel, S., Rjosk, C. & Heinschel, A. (2023). Merkmale der Unterrichtsqualität im Fach Deutsch. In P. Stanat, S. Schipolowski, R. Schneider, S. Weirich, S. Henschel & K. A. Sachse (Hrsg.), *IQB-Bildungstrend 2022. Sprachliche Kompetenzen am Ende der 9. Jahrgangsstufe im dritten Ländervergleich* (S. 359–387). Münster: Waxmann.

- Herdina, P. & Jessner, U. (2002). *A Dynamic Model of Multilingualism: Perspectives of Change in Psycholinguistics*. Clevedon: Multilingual Matters. <https://doi.org/10.21832/9781853595547>
- Herné, K.-L. (2014). Entdecken – verstehen – anwenden. Rechtschreibförderung auf der Grundlage der Aachener förderdiagnostischen Rechtschreibfehler-Analyse. In K. Siekmann (Hrsg.), *Theorie, Empirie und Praxis effektiver Rechtschreibdiagnostik* (Stauffenburg Deutschdidaktik, Bd. 2, 141–154). Tübingen: Stauffenburg Verlag.
- Herné, K.-L. & Naumann, C. L. (2016). *Aachener förderdiagnostische Rechtschreibfehler-Analyse. Systematische Einführung in die Praxis der Fehleranalyse mit Auswertungshilfen zu insgesamt 33 standardisierten Testverfahren* (5., überarbeitete Auflage). Aachen: Alfa Zentaurus.
- Hincha, X. (2003). Ist Hanyu Pinyin eine Schrift? *CHUN – Chinesischunterricht*, (18), 145–153.
- Hinney, G. (1997). *Neubestimmung von Lerninhalten für den Rechtschreibunterricht. Ein fachdidaktischer Beitrag zur Schriftaneignung als Problemlöseprozess*. Frankfurt: Peter Lang.
- Hirschmann, H. (2019). *Korpuslinguistik*. Stuttgart: J. B. Metzler Verlag. <https://doi.org/10.1007/978-3-476-05493-7>
- Hirschmann, H., Binanzer, A. & Langlotz, M. (2023). NaLeKo: Ein komplex annotiertes Lernerkorpus mit schriftlichen Erzähltexten des Deutschen als Erst- und Zweitsprache. *Korpora Deutsch als Fremdsprache*, 3(2), 230–238. <https://doi.org/10.48694/kor-daf.3859>
- Hirschmann, H., Lüdeling, A., Schadrova, A., Bobeck, D., Klotz, M., Akbari, R. et al. (2022). FALKO: Eine Familie vielseitig annotierter Lernerkorpora des Deutschen als Fremdsprache. *Korpora Deutsch als Fremdsprache*, 2(2), 139–148.
- Hiver, P. (2014). Attractor States. In Z. Dörnyei, P. D. MacIntyre & H. Alastair (Hrsg.), *Motivational Dynamics in Language Learning* (pp. 20–28). Tallinn: Multilingual Matters. <https://doi.org/10.13140/RG.2.1.2501.8722>
- Hofmann, N. (2008). *Unterrichtsexpertise und Rechtschreibleistungen*. Heidelberg, Pädag. Hochsch., Diss., 2008. <https://doi.org/75021>
- Höhle, B. (2010). Erstspracherwerb: Wie kommt das Kind zur Sprache? In B. Höhle (Hrsg.), *Psycholinguistik* (S. 125–139). Berlin: Akademie Verlag. <https://doi.org/10.1524/9783050052861.125>
- Holland, J. H. (2000). *Emergence: From Chaos to Order* (Revised ed.). Oxford: Oxford University Press.
- Huang, B. & Liao, X. (2017). *Xiandai hanyu: Modern Chinese*. Beijing: Higher Education Press.
- Hufeisen, B. (2003). L1, L2, L3, L4, Lx – alle gleich? Linguistische, lernerinterne und lernerexterne Faktoren in Modellen zum multiplen Spracherwerb. *Zeitschrift für Interkulturellen Fremdsprachenunterricht*, 8(2), 97–109.
- Hufeisen, B. & Neuner, G. (2005). *Mehrsprachigkeitskonzept. Tertiärsprachenlernen; Deutsch nach Englisch* (2. Dr., korrigierte Aufl.). Strasbourg: Council of Europe Publishing.

- Hulstijn, J. H., Ellis, Rod & Søren, W. E. (2015). Orders and sequences in the acquisition of L2 morphosyntax, 40 years on: An introduction to the special issue. *Language Learning*, 65(1), 1–5. <https://doi.org/10.1111/lang.12097>
- Huneke, H.-W. & Steinig, W. (2013). *Deutsch als Fremdsprache. Eine Einführung* (Grundlagen der Germanistik, Bd. 34, 6., neu bearbeitete und erweiterte Auflage). Berlin: Schmidt. Verfügbar unter: <https://esv-elibrary.de/book/99.160005/9783503200993>
- Hunold, C. (2009). *Untersuchungen zu segmentalen und suprasegmentalen Ausspracheabweichungen chinesischer Deutschlernender*. Dissertation.
- Hunold, C. (2021). Ausgangssprache Chinesisch. In S. Dahmen, U. Hirschfeld, S. Meißner & K. Reinke (Hrsg.), *Kontrastive Phonetik für Deutsch als Fremd- und Zweitsprache* (S. 1–9). Berlin: Erich Schmidt Verlag.
- Hurschler Lichtsteiner, S., Wicki, W. & Falmann, P. (2018). Impact of handwriting training on fluency, spelling, and text quality among third graders. *Reading and Writing*, 31(6), 1295–1318. <https://doi.org/10.1007/s11145-018-9825-x>
- Jacobs, J. (2005). *Spatien. Zum System der Getrennt- und Zusammenschreibung im heutigen Deutsch*. Berlin: De Gruyter. <https://doi.org/10.1515/9783110919295>
- Jaeuthe, J. (2023). *Die Entwicklung der Rechtschreibkompetenz von Grundschulkindern unter Berücksichtigung von individuellen Merkmalen und Merkmalen der Unterrichtsqualität*. Potsdam: Universität Potsdam Verlag. <https://doi.org/10.25932/publishup-61527>
- Jagemann, S. (2015). Was wissen Studierende über die Regeln der deutschen Wortschreibung? – Eine explorative Studie zum graphematischen Wissen. In C. Bräuer & D. Wieser (Hrsg.), *Lehrende im Blick. Empirische Lehrerforschung in der Deutschdidaktik* (S. 255–279). Wiesbaden: Springer VS.
- Jagemann, S. (2016). „Hör mal genau hin: <Tru – He>.“ Wie verstehen und erklären angehende Lehrer/innen das silbeninitiale-h? In H. Zimmermann & A. Peyer (Hrsg.), *Wissen und Normen – Facetten professioneller Kompetenz von Deutschlehrkräften* (Germanistik – Didaktik – Unterricht, Band 16, S. 221–246). Frankfurt am Main: Peter Lang.
- Jeuk, S. (2012). Orthographieerwerb mehrsprachiger Kinder in der ersten Klasse. In W. Grieshaber & Z. Kalkan-Aydın (Hrsg.), *Orthographie- und Schriftspracherwerb bei mehrsprachigen Kindern* (1. Aufl.). Stuttgart: Fillibach Verlag; Klett.
- Jeuk, S. (2021). *Deutsch als Zweitsprache in der Schule. Grundlagen – Diagnose – Förderung* (5. Auflage). Stuttgart: Kohlhammer Verlag. <https://doi.org/10.17433/978-3-17-039003-4>
- Jordan, G. (2004). *Theory Construction in Second Language Acquisition* (Language Learning & Language Teaching, Bd. 8). Amsterdam, Philadelphia: John Benjamins. <https://doi.org/10.1075/llt.8>
- Jung, B. & Günther, H. [Herbert]. (2016). *Erstsprache, Zweitsprache, Fremdsprache. Eine Einführung* (Beltz Pädagogik, Neuausgabe, 3., überarbeitete und erweiterte Aufl.). Weinheim: Beltz. Verfügbar unter: <http://nbn-resolving.org/urn:nbn:de:bsz:31-epflicht-1120420> (13.06.2026)

- Kalkavan-Aydın, Z. (2012). Orthographische Markierungen des Deutschen in türkischsprachigen Lernertexten: Ein Indiz für unzureichende Orthographiekennntnisse in der Erstsprache oder Transferleistungen und Rechtschreibbewusstsein? In W. Grieshaber & Z. Kalkavan-Aydın (Hrsg.), *Orthographie- und Schriftspracherwerb bei mehrsprachigen Kindern* (1. Aufl., S. 57–80). Stuttgart: Fillibach Verlag; Klett.
- Kargiotidis, A. & Manolitsis, G. (2024). Are children with early literacy difficulties at risk for anxiety disorders in late childhood? *Annals of Dyslexia*, 74(1), 82–96. <https://doi.org/10.1007/s11881-023-00291-7>
- Karlgren, B. (2001). *Schrift und Sprache der Chinesen*. Würzburg: Springer-Verlag Berlin Heidelberg. <https://doi.org/10.1007/978-3-642-56465-9>
- Klein, W. (2000). Prozesse des Zweitspracherwerbs. In H. Grimm (Hrsg.), *Enzyklopädie der Psychologie* (Enzyklopädie der Psychologie Themenbereich C Theorie und Forschung, Ser. 3 Sprache; Bd. 3, S. 538–570). Göttingen: Hogrefe Verl. für Psychologie.
- Kleiner, S. (Hrsg.). (2011 ff.). /h/ im Inlaut [AADG-Karte]. In: Atlas zur Aussprache des deutschen Gebrauchsstandards (AADG). Leibniz-Institut für Deutsche Sprache (IDS). <http://prowiki.ids-mannheim.de/bin/view/AADG/HimInlaut> (13.06.2026)
- Kohler, K. J. (1995). *Einführung in die Phonetik des Deutschen* (Grundlagen der Germanistik, Bd. 20, 2., neubearb. Aufl.). Berlin: Erich Schmidt Verlag.
- Krause, T. & Zeldes, A. (2016). ANNIS3: A new architecture for generic corpus query and visualization. *Digital Scholarship in the Humanities*, 31(1), 118–139. <https://doi.org/10.1093/llc/fqu057>
- Krifka, M., Błaszczak, J., Leßmöllmann, A., Meinunger, A., Stiebels, B., Tracy, R. et al. (2014). *Das mehrsprachige Klassenzimmer*. Berlin, Heidelberg: Springer-Verlag. <https://doi.org/10.1007/978-3-642-34315-5>
- Krumm, H.-J., Fandrych, C., Hufeisen, B. & Riemer, C. (2010). *Deutsch als Fremd- und Zweitsprache. Ein internationales Handbuch* (Handbücher zur Sprach- und Kommunikationswissenschaft, 35.1/35.2). Berlin, New York: De Gruyter. <https://doi.org/10.1515/9783110240245>
- Lado, R. (1957). *Linguistics across Cultures: Applied Linguistics and Language Teachers*. Ann Arbor: University of Michigan Press.
- Lado, R. (1961). *Language Testing. The Construction and Use of Foreign Language Tests*. London: Langmans, Green & Co., Ltd.
- Lado, R. (1967). *Moderner Sprachunterricht. Eine Einführung auf wissenschaftlicher Grundlage*. München: Max Niemeyer Verlag.
- Larsen-Freeman, D. (1997). Chaos/complexity science and second language acquisition. *Applied Linguistics*, 18(2), 141–165. <https://doi.org/10.1093/applin/18.2.141>
- Larsen-Freeman, D. (2017). Complexity Theory. The lessons continue. In *Complexity Theory and Language Development. In Celebration of Diane Larsen-Freeman* (S. 11–50). Amsterdam: John Benjamins. <https://doi.org/10.1075/llt.48.02lar>
- Larsen-Freeman, D. & Cameron, L. (2008). Research Methodology on Language Development from a Complex Systems Perspective. *The Modern Language Journal*, 92(2), 200–213. <https://doi.org/10.1111/j.1540-4781.2008.00714.x>

- Lavalaye, A. (2015). „ABER MAN KANN ALLES ERLERNEN, WENN MAN *WILL HAT.“. *Eine Fehlerlinguistische Analyse schriftlicher Textproduktionen von Mazedonischsprachigen Deutschlernern (L3)*. Dissertation. Universität Würzburg, Würzburg.
- Lay, T. (2017). Lernschwierigkeiten chinesischer Deutschlerner im Anfängerunterricht unter Berücksichtigung des Chinesischen als Erstsprache und des Englischen als erster Fremdsprache. *Zielsprache Deutsch*, 44(2), 19–38.
- Lenth, R. V. (2025). Estimated Marginal Means, aka Least-Squares Means (Version [R package emmeans version 1.10.7] [Computer software]). Comprehensive R Archive Network (CRAN). Verfügbar unter: <https://cran.r-project.org/web/packages/emmeans/index.html> (13.06.2026)
- Li, F. (2014). *Lernersprachliche Entwicklung chinesischer Deutschlerner – Eine empirische Untersuchung zum Erwerb der substantivischen Pluralmarkierungen*. Diss. Ruhr-Universität Bochum, Bochum.
- Li, Y. [Yuan] & Ge, N. (2019). Digitalisierung in der chinesischen Germanistik am Beispiel der Korpuslinguistik im digitalen Zeitalter 1. *Jahrbuch für Internationale Germanistik*, LI(2), 191–203. https://doi.org/10.3726/JA512_191
- Li, Y. [Yuan] & Wu, Z. (2022a). Chinesisches Deutschlerner-Korpus (CDLK): Ein umfangreiches Korpus mit Mehrebenen-Annotation und multidimensionalen Metadaten. In M. Kupietz & T. Schmidt (Hrsg.), *Neue Entwicklungen in der Korpuslandschaft der Germanistik. Beiträge zur IDS-Methodenmesse 2022*. (Korpuslinguistik und interdisziplinäre Perspektiven auf Sprache 11, S. 223–236). Tübingen: Narr. <https://doi.org/10.24053/9783823396024>
- Li, Y. [Yuan] & Wu, Z. (2022b, 16. März). *Chinesisches Deutschlerner-Korpus (CDLK). METHODENMESSE*, Online-Konferenz.
- Liang, M. & Nerlich, M. (2009). *Studienweg Deutsch. Kursbuch 1* (Bd. 1). Beijing: Foreign Language teaching and research press.
- Lightbown, P. M. & Spada, N. (2013). *How Languages Are Learned* (4th ed.). Oxford University Press.
- Lin, T. & Wang, L. (2013). *Yu Yin Xue Jiao Cheng. A Course in Phonetics*. Beijing: Beijing University Press.
- Liu, C. (1982). Einige Überlegungen zum gegenwärtigen Deutschunterricht in China. *Zielsprache Deutsch*, (4), 29–39.
- Liu, M. (2019). *Lexikalische Fehleranalyse chinesischer Deutschlernender – eine empirische Untersuchung am Benediktinergymnasium Ettal*. Masterarbeit. Otto-Friedrich-Universität Bamberg, Bamberg.
- Liu, T. (2015). „Ich verstehe nur Chinesisch“ – Kontrastierung der chinesischen und deutschen Phonetik/Phonologie als Basis für die Entwicklung von Lehr- und Lernmaterialien für Deutschlernende chinesischer Muttersprache. Diss. Humboldt-Universität zu Berlin, Berlin.
- Lü, C. (2017). The Roles of Pinyin Skill in English-Chinese Bilingual Learning: Evidence From Chinese Immersion Learners. *Foreign Language Annals*, 50(2), 306–322. <https://doi.org/10.1111/flan.12269>

- Lüdeling, A. & Hirschmann, H. (2015). Error annotation systems. In S. Granger, G. Gilquin & F. Meunier (Hrsg.), *The Cambridge Handbook of Learner Corpus Research* (S. 135–158). Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9781139649414.007>
- Ma, J. (1984). Möglichkeiten, Probleme und Methoden des deutsch-chinesischen Grammatikvergleichs. In H.-R. Fluck (Hrsg.), *Te-han-yü-yen-pi-chiao. Sprachvergleichende Arbeiten in d. Bereichen Phonetik/Phonologie – Lexik/Morphologie/Syntax-Übers. – Didaktik an d. Tongji-Univ. Shanghai = Kontrastive Linguistik deutsch-chinesisch* (S. 22–69). Heidelberg: Groos.
- Maas, U. (2016). Migrationsschwelle Sprachausbau. Ein gemeinsames Projekt mit Michael Bommers. In Vorstand des Instituts für Migrationsforschung und Interkulturelle Studien (IMIS) der Universität Osnabrück (Hrsg.), *IMIS-BEITRÄGE* (Heft 50 (2016), S. 1–174).
- Maaz, K. & Lörz, M. (2024). *20 Jahre PISA: Soziale Bildungsungleichheit im Fokus*. Gütersloh: Bertelsmann Stiftung. <https://doi.org/10.11586/2024112>
- Mannhaupt, G. (2003). Früherkennung und Prävention von Problemen im Schriftspracherwerb. *Grundschule*, 35(9), 45–48.
- Marioulas, J. & Wu, L. (2015). Expansion und Hierarchisierung der chinesischen Germanistik. *German as a Foreign Language*, (3), 30–50.
- Marsh, G., Friedman, M. & Desberg, P. (1981). A cognitive-developmental theory of reading acquisition. In G. E. MacKinnon & T. G. Waller (Hrsg.), *Reading research: Advances in theory and practice* (S. 91–118). London: Academic Press.
- May, P. (1990). Kinder lernen rechtschreiben: Gemeinsamkeiten und Unterschiede guter und schwacher Lerner. In H. Brügelmann & H. Balhorn (Hrsg.), *Das Gehirn, sein Alphabet und andere Geschichten* (S. 245–253). Konstanz: Faude.
- May, P. (2013). *HSP 1–10. Hamburger Schreib-Probe. Manual/Handbuch: Diagnose orthografischer Kompetenz* (Neunormierung 2012, 1. Auflage). Zur Erfassung der grundlegenden Rechtschreibstrategien. Stuttgart: vpm verlag für pädagogische medien.
- Mazur, A. & Quignard, M. (2025). Pauses in the Dynamics of Handwriting Production: Evidence of Persistent Difficulties in French Students With Dyslexia. *Dyslexia (Chichester, England)*, 31(1), e1789. <https://doi.org/10.1002/dys.1789>
- McBride-Chang, C., Bialystok, E., Chong, K. K. Y. & Li, Y. [Yanping]. (2004). Levels of phonological awareness in three cultures. *Journal of Experimental Child Psychology*, 89(2), 93–111. <https://doi.org/10.1016/j.jecp.2004.05.001>
- Meinhold, G. & Sock, E. (1982). *Phonologie der deutschen Gegenwartssprache* (2., durchgesehene Auflage). Leipzig: VEB Bibliographisches Institut Leipzig.
- Meisel, J. M. (2012). *First and Second Language Acquisition*. Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9780511862694>
- Meisel, J. M., Clahsen, H. & Pienemann, M. (1981). On determining developmental stages in natural second language acquisition. *Studies in Second Language Acquisition*, 3(2), 109–135. <https://doi.org/10.1017/S0272263100004137>
- MERLIN-Projekt. (2014a). MERLIN annotation architecture. Verfügbar unter: <https://merlin-platform.eu/docs/Annotation%20guidelines.pdf> (13.06.2026)

- MERLIN-Projekt. (2014b). Nutzerhandbuch. Verfügbar unter: www.merlin-platform.eu
- Mesch, B. (2011). Konzepte des Erwerbs der Getrennt- und Zusammenschreibung. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5, S. 268–295). Baltmannsweiler: Schneider Verlag Hohengehren.
- Müller, A. (2010). *Rechtschreiben lernen. Die Schriftstruktur entdecken; Grundlagen und Übungsvorschläge* (Praxis Deutsch, 1. Aufl.). Seelze, Stuttgart: Kallmeyer in Verb. mit Klett; Klett.
- Müller, H. G. (2008). *Das Onlineportal Orthografietrainer.net – Lernpsychologische und didaktische Überlegungen zu einem automatisierten digitalen Übungskonzept*. Verfügbar unter: <https://www.orthografietrainer.net/service/Psychologisch-didaktischer%20Hintergrund%20Orthografietrainer.net.pdf>
- Müller, H.-G. (2007). *Zum „Komma nach Gefühl“: Implizite und explizite Kommakompetenz von Berliner Schülerinnen und Schülern im Vergleich*. Frankfurt am Main: Peter Lang GmbH, Internationaler Verlag der Wissenschaften.
- Müller, H.-G. & Schroeder, C. (2022). On the influence of the first language on orthographic competences in German as a second language: a comparative analysis. *Applied Linguistics Review*, 15(2), 1–25. <https://doi.org/10.1515/applirev-2020-0145>
- Nevyjel, E. (2020). *Rechtschreibung üben – aber richtig! Übungsbuch zur Wortschatzarbeit und zur Rechtschreibung für Kinder mit DaZ/Förderbedarf ab der 3. Schulstufe* (1. Auflage). Eisenstadt: E. Weber Verlag GmbH. Verfügbar unter: <https://permalink.obvsg.at/AC16189825>
- Niemietz, J., Jindra, C., Schneider, R., Schumann, K., Schipolowski, S. & Sachse, K. A. (2023). Soziale Disparitäten. In P. Stanat, S. Schipolowski, R. Schneider, S. Weirich, S. Henschel & K. A. Sachse (Hrsg.), *IQB-Bildungstrend 2022. Sprachliche Kompetenzen am Ende der 9. Jahrgangsstufe im dritten Ländervergleich* (S. 261–295). Münster: Waxmann.
- Nimz, K. (2022). Orthographische Kompetenzen deutsch-türkisch bilingualer Schüler*innen: Sekundäranalyse des MULTILIT-Korpus. In K. Nimz, C. Noack & K. Schmidt (Hrsg.), *Mehrsprachigkeit und Orthographie. Empirische Studien an der Schnittstelle von Linguistik und Sprachdidaktik* (Thema Sprache – Wissenschaft für den Unterricht, S. 31–49). Bielefeld: wbv Publikation. <https://doi.org/10.3278/9783763973644>
- Nolda, A. (2019). *Annotation von Lernerdaten mit EXMARALDA (Dulko)* (Berlin-Brandenburgische Akademie der Wissenschaften). Verfügbar unter: <http://andreas.nolda.org>
- Nolda, A. (2024). DIE DULKO-TOOLS DES EXMARALDA-PARTITUR-EDITORS. Von einer externen Toolsammlung zum integrierten Bestandteil. *Korpora Deutsch als Fremdsprache*, 4(2), 224–235.
- Offizielle Norm des Ministry of Education der VR China. (2012). Basic rules of Chinese phonetic alphabet orthography. GB/T 16159–2012.
- Ortner, L., Kühnhold, I., Wellmann, H. & Pümpel-Mader, M. (1991). *Deutsche Wortbildung. Typen und Tendenzen in der Gegenwartssprache* (Sprache der Gegenwart, Bd. 79). Düsseldorf: Pädagog. Verl. Schwann.

- Ossner, J. (2006). Kompetenzen und Kompetenzmodelle im Deutschunterricht. *Didaktik Deutsch*, 11(21), 5–19.
- Peperkamp, S. & Dupoux, E. (2002). A typological study of stress 'deafness'. In C. Gussenhoven & N. Warner (Eds.), *Laboratory Phonology 7* (pp. 203–240). Mouton de Gruyter.
- Pießnack, C. & Schübel, A. (2005). Untersuchungen zur orthographischen Kompetenz von Abiturientinnen und Abiturienten im Land Brandenburg. In Zentrum für Lehrerbildung (Hrsg.), *Fachdidaktik* (S. 50–73). Potsdam: Universitätsverlag Potsdam.
- Pompino-Marschall, B. (2003). *Einführung in die Phonetik* (De Gruyter Studienbuch, 2. durchges. und erw. Aufl., Reprint 2020). Berlin, Boston: De Gruyter. <https://doi.org/10.1515/9783110913248>
- Posit team. (2023). RStudio: Integrated Development Environment for R. (Version 420) [Computer software]. Verfügbar unter: URL <http://www.posit.co/>.
- Pracht, H. (2012). *Schemabasierte Basisalphabetisierung im Deutschen. Ein Praxisbuch für Lehrkräfte*. Münster, New York, München, Berlin: Waxmann. Verfügbar unter: <https://elibrary.utb.de/doi/book/10.31244/9783830976233>
- Qian, W. (2006). *Kontrastive Angewandte Studien Chinesisch-Deutsch*. Beijing: Foreign Language teaching and research press.
- Qiu, M. (1984a). Kontrastive Untersuchung der Phonemsysteme des Deutschen und des Chinesischen. In H.-R. Fluck (Hrsg.), *Tè-han-yü-yen-pi-chiao. Sprachvergleichende Arbeiten in d. Bereichen Phonetik/Phonologie – Lexik/Morphologie/Syntax-Übers. – Didaktik an d. Tongji-Univ. Shanghai = Kontrastive Linguistik deutsch-chinesisch* (S. 96–132). Heidelberg: Groos.
- Qiu, M. (1984b). Strukturformel der Phonotagmen im Deutschen und im Chinesischen. In H.-R. Fluck (Hrsg.), *Tè-han-yü-yen-pi-chiao. Sprachvergleichende Arbeiten in d. Bereichen Phonetik/Phonologie – Lexik/Morphologie/Syntax-Übers. – Didaktik an d. Tongji-Univ. Shanghai = Kontrastive Linguistik deutsch-chinesisch* (S. 76–95). Heidelberg: Groos.
- R Core Team. (2022). R: A language and environment for statistical computing. R Foundation for Statistical Computing [Computer software]. Wien.
- Rasool, U., Aslam, M. Z., Mahmood, R., Barzani, S. H. H. & Qian, J. (2023). Pre-service EFL teacher's perceptions of foreign language writing anxiety and some associated factors. *Heliyon*, 9(2), e13405. <https://doi.org/10.1016/j.heliyon.2023.e13405>
- Reichardt, A. (2015). *Rechtschreibung im Textraum – Modellierungen der Schreibkompetenz in der Grundschule* (Kölner Beiträge zur Sprachdidaktik, Bd. 9). Köln: Universitäts- und Stadtbibliothek Köln. <https://doi.org/64417>
- Reichardt, M. & Reichardt, S. (1990). *Grammatik des modernen Chinesisch*. Leipzig: Verlag Enzyklopädie Leipzig.
- Ren, R. (2022). Educational Success in Transitional China: The Gaokao and Learning Capital in Elite Professional Service Firms. *The China Quarterly*, 252, 1277–1298. <https://doi.org/10.1017/S0305741022000856>
- Reznicek, M., Lüdeling, A. & Hirschmann, H. (2013). Competing target hypotheses in the Falko corpus: A flexible multi-layer corpus architecture. In Ana Díaz-Negrillo, N. Ballier & P. Thompson (Hrsg.), *Automatic Treatment and Analysis of Learner Corpus Data* (S. 101–123). Amsterdam: Benjamins. <https://doi.org/10.1075/sci.59.07rez>

- Reznicek, M., Lüdeling, A., Krummes, C., Schwantuschke, F., Walter, M., Schmidt, K. [Karin] et al. (2012). *Das Falko-Handbuch Korpusaufbau und Annotationen. Version 2.01*. Verfügbar unter: <https://www.linguistik.hu-berlin.de/de/institut/professuren/korpuslinguistik/forschung/falko/FalkoHandbuchV2/view>
- Richards, J. C. (1974). A Non-Contrastive Approach to Error Analysis. In J. C. Richards (Ed.), *Error Analysis. Perspectives on Second Language Acquisition*. London: Longman.
- Richter, A. (2008). *Orthographische Kompetenz in Deutsch unter der Bedingung von Zweisprachigkeit. Rechtschreibfertigkeiten italienischer und deutscher Schüler der Klassenstufen 4, 5, 6 und 9*. Diss. Technische Universität Dresden, Dresden.
- Röber, C. (2011). Konzepte des Erwerbs der Groß-/Kleinschreibung. In U. Bredel & T. Reißig (Hrsg.), *Weiterführender Orthographieerwerb* (Deutschunterricht in Theorie und Praxis, Bd. 5, S. 296–317). Baltmannsweiler: Schneider Verlag Hohengehren.
- Röber, C. (2012). Die Orthografie als Lehrmeisterin beim Spracherwerb. Die Bedeutung der Rechtschreibung für die Veranschaulichung der Strukturen des Deutschen im DaZ-Unterricht. *Deutsch als Zweitsprache*, 2012, 34–49.
- Röber, C. (2013). *Die Leistungen der Kinder beim Lesen- und Schreibenlernen. Grundlagen der silbenanalytischen Methode; ein Arbeitsbuch mit Übungsaufgaben*. Bielefeld: wbv Publikation. Verfügbar unter: <https://elibrary.utb.de/doi/book/10.3278/9783763965526>
- Röber-Siekmeyer, C. (1997). *Die Schriftsprache entdecken. Rechtschreiben im offenen Unterricht* (Reihe Werkstattbuch Grundschule, 3., erg. und neu ausgestattete Aufl., [mit Anm. zur neuen Rechtschreibung]). Weinheim, Basel: Beltz.
- Roche, J. (2013). *Fremdsprachenerwerb – Fremdsprachendidaktik* (utb-studi-e-book, Bd. 2691, 3., vollständig überarbeitete Auflage). Tübingen, Basel: Francke Verlag. Verfügbar unter: <http://www.utb-studi-e-book.de/9783838540382>
- Roos, J. & Schöler, H. (2009). *Entwicklung des Schriftspracherwerbs in der Grundschule*. Wiesbaden: VS Verlag für Sozialwissenschaften. <https://doi.org/10.1007/978-3-531-91574-6>
- Rösler, D. (1982). Rechtschreibung und Rechtschreibreform – Kein Thema für Deutsch als Fremdsprache? *Zielsprache Deutsch*, 13, 29–35. Verfügbar unter: <https://jilupub.uni-giessen.de/bitstream/handle/jilupub/16960/Roesler-Rechtschreibung.pdf?sequence=1> (13.06.2026)
- Ruppert, C. & Hanulíková, A. (2022). Die Rechtschreibleistung ein- und mehrsprachiger Schüler*innen: Fehlerraten und Fehlerarten. In K. Nimz, C. Noack & K. Schmidt (Hrsg.), *Mehrsprachigkeit und Orthographie. Empirische Studien an der Schnittstelle von Linguistik und Sprachdidaktik* (Thema Sprache – Wissenschaft für den Unterricht). Bielefeld: wbv Publikation.
- Scheerer-Neumann, G. (1987). Kognitive Prozesse beim Rechtschreiben: eine Entwicklungsstudie. In G. Eberle & G. Reiß (Hrsg.), *Probleme beim Schriftspracherwerb. Möglichkeiten ihrer Vermeidung und Überwindung* (S. 193–219). Heidelberg: Heidelberger Verlagsanstalt und Druckerei GmbH.
- Scheerer-Neumann, G. (2006). Entwicklung der basalen Lesefähigkeit. In U. Bredel, H. Günther, P. Klotz, J. Ossner & G. Siebert-Ott (Hrsg.), *Didaktik der deutschen Sprache* (S. 513–524). Paderborn: Schöningh.

- Schmid, H. (1994). *Probabilistic Part-of-Speech Tagging Using Decision Trees. Proceedings of the International Conference on New Methods in Language Processing*. Manchester. Verfügbar unter: <https://www.cis.lmu.de/~schmid/tools/TreeTagger/>
- Schmidt, K. [Karsten] (2022). Orthographie, Mehrsprachigkeit und Sprachdidaktik: Einleitung in den Sammelband. In K. Nimz, C. Noack & K. Schmidt (Hrsg.), *Mehrsprachigkeit und Orthographie. Empirische Studien an der Schnittstelle von Linguistik und Sprachdidaktik* (Thema Sprache – Wissenschaft für den Unterricht, S. 7–14). Bielefeld: wbv Publikation.
- Schmidt, T. & Wörner, K. (2014). „EXMARaLDA“. In J. Durand, U. Gut & G. Kristoffersen (Hrsg.), *The Oxford Handbook of Corpus Phonology* (S. 402–419). London: Oxford University Press. <https://doi.org/10.1093/oxfordhb/9780199571932.013.030>
- Schmidt, W. R. (1990). The role of consciousness in second language learning. *Applied Linguistics*, 11(2), 129–156. <https://doi.org/10.1093/applin/11.2.129>
- Schneider, H., Becker-Mrotzek, M., Sturm, A., Jambor-Fahlen, S., Neugebauer, U., Efig, C. et al. (2013). *Expertise zur Wirksamkeit von Sprachförderung*. Expertise. Fachhochschule Nordwestschweiz, Pädagogische Hochschule, Köln.
- Schroeder, C. & Schellhardt, C. (2015). *MULTILIT: manual, criteria of transcription and analysis for German, Turkish and English*. Potsdam: Universitätsverlag Potsdam.
- Schwendemann, M. (2023). *Die Entwicklung syntaktischer Strukturen. Eine Längsschnittstudie anhand schriftlicher Sprachdaten erwachsener Deutschlernender mit der Erstsprache Arabisch* (Studien Deutsch als Fremd- und Zweitsprache, Bd. 17, 1st ed. 2023). Berlin: Erich Schmidt Verlag; Imprint Erich Schmidt Verlag GmbH & Co. KG. <https://doi.org/10.37307/b.978-3-503-21223-1>
- Selinker, L. (1972). Interlanguage. *IRAL – International Review of Applied Linguistics in Language Teaching*, 10(1–4). <https://doi.org/10.1515/iral.1972.10.1-4.209>
- Selinker, L. (2014). Interlanguage 40 years on: Three themes from here. In Z. Han & E. Tarone (Eds.), *Interlanguage. Forty years later* (Language Learning & Language Teaching, v. 39, S. 221–246). Amsterdam: John Benjamins. <https://doi.org/10.1075/llt.39.12ch10>
- Selting, M., Auer, P., Barth-Weingarten, D., Bergmann, J., Bergmann, P., Birkner, K. et al. (2009). Gesprächsanalytisches Transkriptionssystem 2 (GAT 2). *Gesprächsforschung – Online-Zeitschrift zur verbalen Interaktion*, (10), 353–402. Verfügbar unter: <http://www.gespraechsforschung-ozs.de/heft2009/px-gat2.pdf> (13.06.2026)
- Shanghai Ranking. (2024). *Best Chinese Universities Ranking*. Verfügbar unter: <https://www.shanghairanking.com/rankings/arwu/2024> (13.06.2026)
- Siekmann, K. (2021). Wortschatz, Phonem-Graphem-Relationen und Fehlschreibungen bei dem Phonem /f/ in freien Schülertexten der Jahrgangsstufen 3–5. *Lernen und Lernstörungen*, 10(3), 161–167. <https://doi.org/10.1024/2235-0977/a000346>
- Siekmann, K. & Thomé, G. (2018). *Der orthographische Fehler. Grundzüge der orthographischen Fehlerforschung und aktuelle Entwicklungen* (2., aktualisierte Auflage). Oldenburg: isb Institut für sprachliche Bildung – Verlag. Verfügbar unter: http://www.isb-oldenburg.de/pdf/lprob_ortho_fehler.pdf

- Spiekermann, H. (1997). „Syntaxfehler“ von Chinesen in der gesprochenen Fremdsprache Deutsch. In P. Schlobinski (Hrsg.), *Syntax des gesprochenen Deutsch* (S. 263–280). Oppladen: VS Verlag für Sozialwissenschaften.
- Spoelman, M. & Verspoor, M. (2010). Dynamic Patterns in Development of Accuracy and Complexity: A Longitudinal Case Study in the Acquisition of Finnish. *Applied Linguistics*, 31(4), 532–552. <https://doi.org/10.1093/applin/amq001>
- Sprenger, A. (1974). Deutschunterricht in Taiwan. *Zielsprache Deutsch*, (4), 182–190.
- Stanat, P., Schipolowski, S., Schneider, R., Weirich, S., Henschel, S. & Sachse, K. A. (Hrsg.). (2023). *IQB-Bildungstrend 2022. Sprachliche Kompetenzen am Ende der 9. Jahrgangsstufe im dritten Ländervergleich*. Münster: Waxmann. <https://doi.org/10.31244/9783830997771>
- Steinig, W., Betzel, D., Geider, F. J. & Herbold, A. (2009). *Schreiben von Kindern im diachronen Vergleich. Texte von Viertklässlern aus den Jahren 1972 und 2002* (1. Aufl.). München: Waxmann.
- Tang, L. (2003). *Lernersprachliche Abweichungen in Aufsätzen chinesischer Deutschlerner. Eine fehlerlinguistische Studie*. Diss. Osnabrück: Der Andere Verlag.
- Thelen, T. (2010). *Automatische Analyse der orthographischen Automatische Analyse orthographischer Leistungen von Schreibanfängern*. Diss. Osnabrück.
- Thomé, G. (1987). *Rechtschreibfehler türkischer und deutscher Schüler* (Sammlung Groos, Bd. 33). Zugl.: Berlin, Freie Univ., Diss., 1987. Heidelberg: Groos.
- Thomé, G. (1999). *Orthographieverwerb. Qualitative Fehleranalysen zum Aufbau der orthographischen Kompetenz* (Theorie und Vermittlung der Sprache, Bd. 29). Frankfurt am Main: Peter Lang GmbH, Internationaler Verlag der Wissenschaften.
- Thomé, G. (2006). Entwicklung der basalen Rechtschreibkenntnisse. In U. Bredel, H. Günther, P. Klotz, J. Ossner & G. Siebert-Ott (Hrsg.), *Didaktik der deutschen Sprache* (S. 369–379). Paderborn: Schöningh.
- Thomé, G. & Gomolka, J. (2007). Rechtschreiben. In B. Beck & E. Klieme (Hrsg.), *Sprachliche Kompetenzen. Konzepte und Messung*. (S. 140–146). Weinheim: Beltz. <https://doi.org/10.25656/01:3241>
- Thomé, G., Siekmann, K. & Thomé, D. (2016). Phonem-Graphem-Verhältnisse in der deutschen Orthographie: Ergebnisse einer neuen 100.000er-Auszählung. In R. Hofmann & M. Kalmár (Hrsg.), *LRS: Lesen – Rechnen – Schreiben. Ein Handbuch*. (S. 103–116). Wien: Lernen mit Pfiff.
- Thomé, G. & Thomé, D. (2020). *OLFA 3–9. Oldenburger Fehleranalyse für die Klassen 3–9: Instrument und Handbuch zur Ermittlung der orthographischen Kompetenz und Leistung aus freien Texten für die Planung von Fördermaßnahmen: mit farbiger Markierung der orthographischen Entwicklungsphasen: mit einer separaten OLFA-Liste für die Schweiz: mit vielen neuen Übungen* (5., verbesserte Auflage). Oldenburg: isb.
- Valtin, R. (1994). Ein Stufenmodell des Rechtschreiblernens. In I. Naegle (Hrsg.), *Rechtschreibunterricht in den Klassen 1–6. Grundlagen – Erfahrungen – Materialien*. (S. 32–37). Frankfurt am Main: Arbeitskreis Grundschule e. V.

- Van Dijk, M., Verspoor, M. & Lowie, W. (2011). Variability and DST. In M. Verspoor, K. de Bot & W. Lowie (Hrsg.), *A dynamic approach to second language development methods and techniques* (Language Learning & Language Teaching (LL<), vol. 29, S. orthoAmsterdam: John Benjamins. <https://doi.org/10.1075/llt.29.04van>
- Van Geert, P. (1991). A dynamic systems model of cognitive and language growth. *Psychological Review*, 98(1), 3–53.
- Verspoor, M. & Lowie, W. (2021). Complex Dynamic Systems Theory (CDST). In N. Tracy-Ventura & M. Paquot (Hrsg.), *The Routledge Handbook of Second Language Acquisition and Corpora* (S. 189–200). London: Routledge, Taylor & Francis Group. <https://doi.org/10.4324/9781351137904-17>
- Verspoor, M., Lowie, W. & Wieling, M. (2021). L2 Developmental Measures from a Dynamic Perspective. In B. Le Bruyn & M. Paquot (Hrsg.), *Learner Corpora Research Meets Second Language Acquisition* (S. 172–190). London: Cambridge University Press. <https://doi.org/10.1017/9781108674577.009>
- Voss, A., Blatt, I. & Kowalski, K. (2007). Zur Erfassung orthografischer Kompetenz in IGLU 2006: dargestellt an einem sprachsystematischen Test auf Grundlage von Daten aus der IGLU-Voruntersuchung. *Didaktik Deutsch*, (23), 15–32.
- Wachendorf, P. (2022). *Rechtschreiben 1 (DaZ). Das Selbstlernheft*. Brühl: Jandorfverlag KG.
- Wang, J. (2017). *Lexikalische Interferenzen zwischen Deutsch und Chinesisch. Eine didaktische Reflexion in Bezug auf das lexikalische Lernen im chinesischen Deutschunterricht*. Marburg: Philipps-Universität Marburg. Verfügbar unter: <http://archiv.ub.uni-marburg.de/diss/z2017/0540/pdf/djw.pdf> (13.06.2026)
- Wang, M. (1988). Hauptprobleme der Deutschlernenden aus China mit der deutschen Aussprache. *Informationen Deutsch als Fremdsprache*, 15(1), 76–82. <https://doi.org/10.1515/infodaf-1988-150112>
- Wang, M. (1993). *Kontrastierung der Phonemsysteme des Chinesischen und des Deutschen*. Frankfurt am Main: Peter Lang.
- Wang, M. (2010). *Chinesisch-deutsche Verständigung: Wo drückt der Schuh? Sprachvergleich und interkulturelle Kommunikation*. Marburg: Tectum.
- Wegner, H. (2003). Entstehung und Funktion der Fugenelemente im Deutschen oder: warum wir keine *Autobahn haben. *Linguistische Berichte*, (196), 425–457. https://doi.org/10.46771/9783967696943_2
- Weingarten, R. (2004). Die Silbe im Schreibprozess und im Schriftspracherwerb. In U. Bredel (Hrsg.), *Schriftspracherwerb und Orthographie* (Diskussionsforum Deutsch, Bd. 16, S. 6–21). Baltmannsweiler: Schneider-Verlag Hohengehren.
- Wind, A. M. & Harding, L. (2020). Attractor States in the Development of Linguistic Complexity in Second Language Writing and the Role of Self-Regulation. In W. Lowie, M. Michel, A. Rousse-Malpat, M. Keijzer & R. Steinkrauss (Hrsg.), *Usage-Based Dynamics in Second Language Development* (S. 130–154). Tallinn: Multilingual Matters. <https://doi.org/10.2307/jj.22730678.11>
- Winter, B. (2020). *Statistics for linguists. An introduction using R*. New York, London: Routledge, Taylor & Francis Group.

- Wiprächtiger-Geppert, M., Riegler, S. & Freivogel, J. (2015). Erfassung des professionellen Wissens von Deutschlehrkräften zu Orthographie und Orthographieverwerb – Forschungsstand und Perspektiven. In C. Bräuer & D. Wieser (Hrsg.), *Lehrende im Blick. Empirische Lehrerforschung in der Deutschdidaktik* (S. 281–300). Wiesbaden: Springer VS. https://doi.org/10.1007/978-3-658-09734-9_13
- Wu, Q. & Borhan, M. T. (2024). Sustainable development of Chinese higher education through comparison of higher education indices. *Frontiers in Education*, 9. <https://doi.org/10.3389/educ.2024.1340637>
- Wu, Z. & Li, Y. [Yuan]. (2022). Entwicklung der lexikalischen Fehler chinesischer Deutschlerner unter dem Einfluss des Chinesischen. 23(1), 121–143. Literaturstraße.
- Xu, Y., Liu, J. & Yu, L. (2016). Analysis of Dictation Errors in Reformed PGG Tests and Implications 德语专四考试改革后的听写错误分析及其启示. *Journal of Ningbo University of Technology*, 28(2), 55–60.
- Yen, C. (1992). *Kontrastive Untersuchungen zur segmentalen Phonetik und Phonologie des Chinesischen und Deutschen*. Erlangen: Palm und Enke.
- Yin, L., Li, W. [Wenling], Chen, X., Anderson, R. C., Zhang, J., Shu, H. et al. (2011). The role of tone awareness and pinyin knowledge in Chinese reading. *Writing Systems Research*, 3(1), 59–68. <https://doi.org/10.1093/wsr/wsr010>
- Zeileis, A., Kleiber, C. & Jackman, S. (2008). Regression Models for Count Data in R. *Journal of Statistical Software*, 27(8), 1–25. <https://doi.org/10.18637/jss.v027.i08>
- Zhao, J. (2020). Deutsch als Fremdsprache in China – aktuelle Situation, Herausforderungen und Ausblick. *Jahrbuch für Internationale Germanistik*, 52(1), 51–64. https://doi.org/10.3726/JA521_51
- Zheng, Y. & Cheng, L. (2018). How does anxiety influence language performance? From the perspectives of foreign language classroom anxiety and cognitive test anxiety. *Language Testing in Asia*, 8(1). <https://doi.org/10.1186/s40468-018-0065-4>
- Zhou, Y. & McBride, C. (2022). Invented spelling in English and pinyin in multilingual L1 and L2 Cantonese Chinese-speaking children in Hong Kong. *Frontiers in Psychology*, 13, 1.039.461.
- Zhu, J. & Best, K.-H. (1992). Zum Wort im modernen Chinesisch. *Oriens Extremus*, (1/2), 45–60. Verfügbar unter: <https://www.jstor.org/stable/24047220> (13.06.2026)

8 Anhang

8.1 Informationen zu den ausgewählten Hochschulen des DeChiLKo-Prüfungskorpus

Kategorie	PGG	Stadt	Provinz	Region	Hochschule
A	2017/2019	Beijing	Beijing	Nordchina (N)	Beijing Institute of Technology (BIT)
A	2017/2019	Wuhan	Hubei	Zentral- und Südchina (ZS)	Wuhan University
A	2017/2019	Guangzhou	Guangdong	Zentral- und Südchina (ZS)	Guangdong University of Foreign Studies
A	2017/2019	Shanghai	Shanghai	Ostchina (O)	East China Normal University (ECNU)
B	2017/2019	Tianjin	Hebei	Nordchina (N)	Tianjin Foreign Studies University (FSU)
B	2017/2019	Jinan	Shandong	Ostchina (O)	Shandong University of Finance and Economics (SDUFE)
B	2017/2019	Mianyang	Sichuan	Südwestchina (SW)	Mianyang Teachers' College
B	2017/2019	Yanbian	Jilin	Nordostchina (NO)	Yanbian University
C	2017/2019	Beijing	Beijing	Nordchina (N)	Beijing City University
C	2017/2019	Xi'an	Shanxi	Nordwestchina (NW)	Northwestern Polytechnical University Ming de College (NPUMD)
C	2017/2019	Shanghai	Shanghai	Ostchina (O)	Shanghai Normal University TIANHUA College (STHU)
C	2017/2019	Chengdu	Sichuan	Südwestchina (SW)	Chengdu Institute of Sichuan International Studies University (CISISU)
D	2017/2019	Tianjin	Hebei	Nordchina (N)	Binhai School of Foreign Affairs of Tianjin Foreign Studies University
D	2019	Changsha	Hunan	Zentral- und Südchina (ZS)	Xingxiang College of Xiangtan University
D	2017	Guangzhou	Guangdong	Zentral- und Südchina (ZS)	Guangdong University of Foreign Studies South China Business College (GWNG)

Kategorie	PGG	Stadt	Provinz	Region	Hochschule
D	2017/2019	Xi'an	Shanxi	Nordwestchina (NW)	Xi'an Fanyi University (XFU)
D	2017/2019	Siping	Jilin	Nordostchina (NO)	Boda College of Jilin Normal University
E	2017/2019	Shenzhen	Guangdong	Zentral- und Südchina (ZS)	Shenzhen Polytechnic (SZPT)
E	2017	Taizhou	Jiangsu	Ostchina (O)	Taizhou University
E	2017/2019	Weihai	Shandong	Ostchina (O)	Shandong Vocational University of Foreign Affairs (SVUFA)
E	2017/2019	Jinan	Shandong	Ostchina (O)	Shandong Youth University of Political Science
E	2019	Chengdu	Sichuan	Südwestchina (SW)	Chengdu Institute of Sichuan International Studies University – Fachrichtung Sprachen und Dolmetschen (CISISU-SD)

8.2 Transkripte der Diktate im DeChiLko-Prüfungskorpus

8.2.1 Transkripte der Diktat-Audiodatei von der PGG-2017

GAT2_Transkript

Durch verschiedene ARten von beLOHnungen (.) WERden die kinder zum WEIterlernen motiviert punkt (.) compUTer sind auch geDULdiger als Lehrer komma (.) wenn man etwas nicht gleich kann punkt (.) AUßerdem lassen sich verSCHIEdene LEHrinhalte (.) DURch die TECHnischen möglichKEiten GANZ deutlich verANSchau-lichen punkt (.) AUS DIEsen gründen sollte man den compUTer (.) nicht von ANfang an Ablehnen punkt (.) viele eltern haben zwar die sorge komma (.) das kinder vor dem computer vereinSAMen komma (.) ABER das kann man verHINdern komma (.) in-DEM man die SPIEL und LERNdauer (.) VORher genau festlegt punkt

IPA_Transkript

'dovɛʃ fɛɐ̯ ʃiːdɒnə ʔaːrtɪn ʔɒn bəˈloːnoʊən ˌveːrdən diː ˈkɪndə ʔtʂom ˈvaɪtəˈlɛənən
motiˈviːɐ̯tə ˈpʊŋktʰ (.) kɔm ˈpjuːtə zɪnt ʔaʊx gəˈdʊlˈdɪgə ʔals ˈleːxə ˈkoma (.) vɛnˈman
ˈɛtvas nɪçt glaiç kan ˈpʊŋktʰ (.) ʔaʊsəˈdeːm ˈlasn̩ zɪç fɛɐ̯ ʃiːdɒnə leːɐ̯ ˈmɪhaltə (.) 'dovɛʃ
diː ˈtʰɛçnɪʃn̩ ˈmøːklɪçkɪɪtɪn ˈgants ˈdɔɪtlɪç fɛɐ̯ ʔanʃaʊlɪçn̩ ˈpʊŋktʰ (.) ʔaʊs ˈdiːzn̩ ˈgɾɪvndɪ
ˈzoltə ˈman deːn kɔm ˈpjuːtə (.) nɪçt ʔɒn ʔanˈfaɪ ʔan ˈʔapˌleːnən ˈpʊŋktʰ (.) ˈfiːlə ˈʔɛltən
ˈhaːbn̩ tʂvaːɐ̯ diː ˈzɔrgə ˈkoma (.) das ˈkɪndə foːɐ̯ deːm kɔn ˈpjuːtə fɛɐ̯ ʔaɪnzə mən
ˈkoma (.) ʔaːbɐ ˈdas kan ˈman fɛɐ̯ ˈhɪndən ˈkoma (.) mˈdeːman diː ˈʃpiːl ʊnt lɛən
ˈdaʊv (.) ˈfoːɐ̯ ˈheːɐ̯ gəˈnaʊ fɛsˌleːkt ˈpʊŋktʰ

8.2.2 Transkripte der Diktat-Audiodatei von der PGG-2019

GAT2_Transkript

ACHTundneuzig PROzent (.) der BUNdesbürger (.) sehen REgelmäßig fern punkt (.) sie WOLlen sich am ABEND (.) vor dem FERNseher unterHALten (.) UND inforMIEren punkt (.) der ALLtag ist STRESSsig komma (.) die LEUte FREUen sich auf das WOCHENende UND wollen sich AUSruhen komma (.) AUSSchlafen (.) ODER mit der (.) FAmilie treffen punkt (.) viele WÜNschen sich mehr ZEIT für HObbys komma (.) SPORT und FREUNDe Punkt (.) auch die Arbeit im GARTen ist beLIEBT und HILFT gegen STRESS punkt (.) EInerseits gibt es MEHR FREIzeitangebote komma (.) ANdererseits müssen die menschen aber SPAREn punkt IMmer mehr deutsche gehen ins SCHWIMMbad komma (.) FAHren FAHRrad und kochen zusammen punkt FREIzeitvergnügen muss nicht IMmer GELD kosten punkt

IPA_Transkript

'axtont,nɔmtʃɪç pʁo'tɛ ɛntʰ (.) de:ɐ 'bʊndəs,bʏkɐʁ (.) 'ze:ənʔ 're:gl̩mɛ:siç 'fɛʁn
'pʊŋktʰ (.) zi: 'vɔlən zɪç ʔam 'ʔa:bəntʰ (.) fo:ɐ de:m 'fɛʁn,ze:ʁ ʔontə'haltənʔ (.) ʔont
ʔɪnfɔʁ'mi:ʁənʔ 'pʊŋktʰ (.) de:ɐ 'ʔal,tak ist 'ʃtɛsiç koma di: 'lɔʏtəʔ 'fɛʁvən zɪç auf das
'vɔxən,ʔendə ʔontʔ 'vɔlən zɪç 'ʔaus,ʁu:ən koma (.) 'ʔaus,ʃla:fən (.) 'ʔo:də mit de:ɐ (.) fa
'mi:liə 'tʁɛfən 'pʊŋktʰ (.) 'fi:lə 'vʏŋʃən zɪç me:ɐ 'tʃaɪtʰ (.) fy:ɐ 'hɔbi:s koma (.) 'ʃpɔktʰ ontʔ
'fɛʁvɛndə 'pʊŋktʰ (.) ax di: 'arbaɪt ɪm 'gartən ist bə'li:ptʰ ʔont 'hɪlftʰ ge:gən 'ʃtɛs 'pʊŋktʰ
(.) ʔamɐ,zarts̩ gi:pt ɛs 'me:ɐ 'fɛʁtsaɪt ʔangəbo:tə koma (.) 'ʔandəʁɐ,zarts̩ mʏsən di:
'mɛnʃən 'ʔa:bə 'ʃpa:rən 'pʊŋktʰ (.) ʔɪmɐ me:ɐ 'dɔʏtʃəʔ 'ge:ən ʔɪns 'ʃvɪm,bat koma (.)
'fa:rən 'fa:ɐ,rat ʔontʔ 'kɔxən tsu'zamən 'pʊŋktʰ (.) 'fɛʁtsaɪt,fɛʁgny:gən mʏs nɪçt 'ʔɪmɐ
'gɛltʰ 'kɔstən 'pʊŋktʰ

8.3 Verfügbare Metadaten im DeChiLKo-Korpus

Metadaten	Variablen	Werte und Erläuterung
Zum Korpus	corpus_title	Deutsch-chinesisches Lernerkorpus Titel des Korpus
	corpus_acronym	DeChiLKo Akronym des Korpus
	distributor	[AUTOR:IN] Einrichtung des Korpus
	availability	free of charge Zugang zum Korpus
	license	CLARIN PUB+BY+SA+PRIV Verfügbarkeit des Korpus
	character_encoding	UTF-8 Zeichenkodierung

Metadaten	Variablen	Werte und Erläuterung
	markup_language	XML Auszeichnungssprache
	corpus_mode	written Modalität des Korpus
	corpus_size	8.219 tokens Anzahl der Token
Zum Text	year	2019 Erhebungsjahr des Texts
	country	China Land des Erstellungsprojekts
	institution	Wuhan University Erhebungsinstitution des Texts
	institution_category	A Typ der Erhebungsinstitution
	region	Zentral- und Südchina Erhebungsregion
	officialLanguageTesting	PGG Sprachtest
	officialLanguageTesting_Type	schriftlich Sprachtesttyp
	text_titel	ein Besuch Titel des Texts
	task_type	Unterrichtsübung Aufgabentyp
	task_instructions	Dreimaliges Vorlesen, Satzzeichen wurden vorgelesen
	written_task	Diktat Aufgabenart
	written_ref_tools	no Durften externe Hilfsmittel verwendet werden?
	text_size	112 tokens Anzahl der Token
	languages used	deu Sprache des Texts
	transcription_guidelines	DeChiLko Handbuch Transkriptionsrichtlinien
	annotation	yes Wurden die Texte annotiert?
	pos_tagged	yes POS-Tagging des Korpus

Metadaten	Variablen	Werte und Erläuterung
	pos_tagset	STTS POS-Tagest des Korpus
	error_annotated	yes Sind die Texte fehlerannotiert?
	error_annotation_tool	EXMARaLDA (Dulko) Tool der Fehlerannotation
	annotation_other	lemmata, sentence spans, target hypothes, differences, annotations about orthographic errors in each position Annotationsebenen
Zu den Autor*innen	age	19 Alter zum Zeitpunkt der Erhebung
	gender	male/female Geschlecht des Lernaler laut Selbstauskunft
	L1	chi Muttersprache
	L2	deu Zweite Sprache
	L2_self-assessment	23/65/75/85 Wie gut beherrscht der Lerner/die Lernerin Deutsch (=L2) nach eigenen Angaben (Schreiben/Lesen/Sprechen/Hören)
	L2_eng	yes Beherrscht der Lerner/die Lernerin Englisch?
	L2_eng_Gaokao	145 Punktestand der Hochschulaufnahmeprüfung (Gaokao), Gesamtpunktzahl ist 150
	L2_eng_CET4	540 Punktestand von CET 4 (The College English Test)
	L2_eng_CET6	550 Punktestand von CET6
	L2_other	Kor/no weitere Fremdsprache
	test_ID	192020006 ID der/des Studierenden in der Prüfung
	total_score	85 Gesamtnote
	score_of_the_written_task	10 Bewertung der Schreibaufgabe
	score_of_the_dictation_task	3 Bewertung der Diktataufgabe
	learner_level	undergraduated/graduated Niveau der Autorin oder des Autors

Metadaten	Variablen	Werte und Erläuterung
	learner_status	<i>student</i> Status der/des Lerner*in
	study_area	<i>German studies</i> Studienfach
	L2_study_monats	1,3 Lerndauer der L2 in Monaten zum Zeitpunkt der Erhebung
	L2_study_institution	<i>University</i> L2-Institution

8.4 Studienverlaufsplan der Germanistikausbildung an der Anhui-Universität

Das Studienjahr		Die Fächer	SWS
1.	1. Semester	- Grundstufe Deutsch I - Hörverstehen Deutsch - Integriertes Fertigkeitstraining Deutsch (Wahlkurs) - andere Fächer (Computerprogrammieren, Sport usw.)	6 2 4 13
	2. Semester	- Grundstufe Deutsch I - Hörverstehen Deutsch - Mündlicher Ausdruck Deutsch - Integriertes Fertigkeitstraining Deutsch (Wahlkurs) - Deutsche Geschichte (Wahlkurs) - andere Fächer (Computerprogrammieren, Sport usw.)	6 2 1 4 2 8
2.	1. Semester	- Grundstufe Deutsch II - Mündlicher Ausdruck Deutsch - Schriftlicher Ausdruck Deutsch - Kritisches Lesen Deutsch - Integriertes Fertigkeitstraining Deutsch (Wahlkurs) - Hörverstehen Deutsch (Mittelstufe) (Wahlkurs) - Deutsche Geschichte (Wahlkurs) - Englisch als zweite Fremdsprache - andere Fächer (Sport, Politik usw.)	4 1 1 2 4 4 2 4 14
	2. Semester	- Grundstufe Deutsch II - Mündlicher Ausdruck Deutsch - Schriftlicher Ausdruck Deutsch - Kritisches Lesen Deutsch - Integriertes Fertigkeitstraining Deutsch (Wahlkurs) - Hörverstehen Deutsch (Mittelstufe) (Wahlkurs) - Deutsche Grammatik (Wahlkurs) - Englisch als zweite Fremdsprache - andere Fächer (Sport, Politik usw.)	4 1 1 2 4 4 2 4 2

Das Studienjahr		Die Fächer	SWS
3.	1. Semester	• Fortgeschrittenes Deutsch • Übersetzen • Landeskunde deutschsprachiger Länder • Deutsche Literaturwissenschaft • Seh- und Hörverstehen mit Sprechtraining (Wahlkurs) • Deutsche Kulturgeschichte (Wahlkurs) • Vermittlung der chinesischen Kultur auf Deutsch (Wahlkurs) • Chinesisch-deutsche Beziehungen (Wahlkurs) • Dolmetschen (Wahlkurs) • Lexikologie (Wahlkurs) • Sprache und Gesellschaft (Wahlkurs) • Einführung in die Ökonomie Deutschlands (Wahlkurs) • Wirtschaftsdeutsch (Wahlkurs)	6 1 2 2 2 2 2 2 2 2 2 2 2
	2. Semester	• Fortgeschrittenes Deutsch • Deutsche Literaturwissenschaft • Deutsche Sprachwissenschaft • Übersetzen • Kulturvergleich Deutschland–China (Wahlkurs) • Interkulturelle Kommunikation Deutschland–China (Wahlkurs) • Einführung in die Außenpolitik Deutschlands (Wahlkurs) • Aktuelle deutsche Presstexte: Lektüre und Analyse (Wahlkurs) • Dolmetschen (Wahlkurs) • Deutsche Stilistik und Rhetorik (Wahlkurs) • wissenschaftliches und technisches Deutsch (Wahlkurs) • Übersetzung und Aufführung deutscher Theaterstücke (Wahlkurs)	6 2 2 2 2 2 2 2 2 2 2 2 2
4.	1. Semester	• Fortgeschrittenes Deutsch • Wissenschaftliches Schreiben • Lektürekurs: Deutsche Romane (Wahlkurs)	6 1 2
	2. Semester	Abschlussarbeit	

Wie lernen chinesische Studierende die deutsche Rechtschreibung und welche systematischen Schwierigkeiten treten dabei auf? Ming Liu untersucht in ihrer Dissertation die Orthographiekompetenz chinesischer Deutschlernender und leistet damit einen Beitrag zur Forschung im Bereich Deutsch als Fremdsprache.

Auf Grundlage des Deutsch-Chinesischen Lernerkorpus DeChiLKo mit über 300 Diktattexten werden Fehlermuster auf phonographischer, silbischer, morphologischer und syntaktischer Ebene analysiert. Der theoretische Rahmen verbindet die Multi-Kompetenz-Hypothese mit der Complex Dynamic Systems Theory und modelliert den Erwerb als nicht-linearen, dynamischen Prozess im Zusammenspiel von L1 Chinesisch, L2 Englisch und L3 Deutsch. Die Ergebnisse zeigen, dass interlinguale Transferphänomene, typologische Distanz und sozioökonomische Bildungsdisparitäten die Rechtschreibleistung beeinflussen. Darauf aufbauend werden didaktische Empfehlungen für den DaF-Unterricht abgeleitet.

Die Arbeit richtet sich an Forschende in der Linguistik und im Bereich Deutsch als Fremdsprache, an Lehrkräfte im chinesischen Germanistikstudium sowie an alle, die sich mit kontrastiver Schriftspracherwerbsforschung befassen.

Ein Schneider Verlag-Titel bei wbv Publikation

wbv



ISBN: 978-3-7639-7936-3

wbv.de