

Allgemeiner Beitrag

Christian Fandrych*, Matthias Schwendemann* und
Franziska Wallner*

„Ich brauch da dringend ein passendes Beispiel ...“: Sprachdidaktisch orientierte Zugriffsmöglichkeiten auf Korpora der ge- sprochenen Sprache aus dem Projekt ZuMult „I urgently need a good example ...“: Access tools facilitating the use of spoken language corpora for language teaching purposes in the context of the ZuMult project

<https://doi.org/10.1515/infodaf-2021-0077>

Zusammenfassung: Im vorliegenden Beitrag werden die im Projekt ZuMult – „Zugänge zu multimodalen Korpora gesprochener Sprache“; Leibniz-Institut für Deutsche Sprache (IDS) Mannheim / Herder-Institut, Univ. Leipzig / Hamburger Zentrum für Sprachkorpora (HZSK), Univ. Hamburg – entwickelten Zugangswege zu Korpora der gesprochenen Sprache vorgestellt. Es handelt sich dabei um digitale Anwendungen, die gezielt für Sprachdidaktiker/-innen geschaffen wurden, um ihnen einen möglichst bedarfsgerechten Zugriff auf authentische gesprochensprachliche Daten zu ermöglichen. Die neu geschaffenen Zugriffsmöglichkeiten sind besonders auf Bedürfnisse der Sprachvermittlung zugeschnitten. So können nunmehr Gesprächsbeispiele aus den mündlichen Korpora FOLK und GWSS¹ an-

1 FOLK ist das Forschungs- und Lehrkorpus Gesprochenes Deutsch und beinhaltet hauptsächlich private, institutionelle und öffentliche Sprechereignisse. GWSS steht für Gesprochene Wissenschaftssprache und umfasst universitäre Prüfungsgespräche sowie studentische und Expertenvorträge im akademischen Kontext. Sowohl FOLK als auch GWSS sind als Korpora in das Korpusmanagementsystem „Datenbank für Gesprochenes Deutsch“ (DGD) des Leibniz-Instituts für Deutsche Sprache in Mannheim integriert (vgl. <https://dgd.ids-mannheim.de>).

***Kontaktpersonen:** Christian Fandrych, E-Mail: fandrych@uni-leipzig.de
Matthias Schwendemann, E-Mail: matthias.schwendemann@uni-leipzig.de
Franziska Wallner, E-Mail: f.wallner@rz.uni-leipzig.de

hand schwierigkeitsbezogener Parameter (wie etwa Wortschatzniveau, Standardnähe/-ferne, Sprechgeschwindigkeit, Anteil typisch mündlicher Phänomene) für die Thematisierung im Unterricht ausgewählt werden. Zudem werden an sprachdidaktischen Nutzungsszenarien orientierte Arbeitsmöglichkeiten mit dem jeweiligen Einzeltranskript angeboten. Der Beitrag zeigt anhand eines sprachdidaktischen Anwendungsszenarios konkret auf, wie die in ZuMult entwickelten digitalen Anwendungen genutzt werden können.

Schlüsselwörter: Korpora gesprochener Sprache, Mündlichkeitsdidaktik, Korpora in der Sprachdidaktik

Abstract: This article presents the access tools for spoken language corpora developed in the ZuMult project – “Access to multimodal spoken language corpora”; Leibniz Institute for the German Language (IDS) Mannheim / Herder Institute, Univ. Leipzig / The Hamburg Center for Language Corpora (HZSK), Univ. Hamburg. These are digital applications that were created specifically for language teaching purposes in order to provide users with access to authentic spoken language data which reflects their professional needs. The newly created tools are particularly tailored to language pedagogy. Thus, examples of conversations from the spoken corpora FOLK and GWSS can now be selected on the basis of difficulty-related parameters (such as vocabulary level, standard proximity/distance, speed of articulation, proportion of typical spoken language phenomena) for discussion in class. In addition, possibilities for working with the respective individual transcript are offered, oriented towards language teaching scenarios. The article demonstrates how the digital applications developed in ZuMult can be used based on a typical language teaching scenario.

Keywords: Spoken language corpora, teaching spoken language skills, language corpora and language pedagogy

1 Sprachdidaktisch orientierte Zugriffsmöglichkeiten auf Korpora der gesprochenen Sprache aus dem Projekt ZuMult

Die Thematisierung der gesprochenen Sprache im DaF/DaZ-Unterricht gilt immer noch als unzureichend. Obwohl Lehr- und Lernmaterialien der höheren Niveaustufen einige typisch mündliche Phänomene wie etwa Modalpartikeln, Interjektionen oder gefüllte Pausen enthalten (vgl. Günthner/Wegner/Weidner 2013),

wird gesprochene Sprache in Hörbeispielen sowie bei der Vermittlung von Redemitteln, Strukturen und kommunikativer Kompetenz immer noch stark vernachlässigt (vgl. Merklin/Riedel 2016; Dietz 2021). Dies liegt nicht nur, aber auch daran, dass didaktisch passende Beispiele authentischer mündlicher Kommunikation nur schwer zu finden und in geeigneter Weise in den Unterricht einzubinden sind.

Korpora der gesprochenen Sprache eröffnen Lehrenden und Lernenden zwar prinzipiell die Möglichkeit, sich situations- und adressatenangemessen mit authentischer mündlicher Kommunikation auseinanderzusetzen; allerdings orientieren sich die entsprechenden Zugriffsmöglichkeiten bislang vor allem an Kategorien, die für sprachwissenschaftliche Analysen Relevanz besitzen. Für sprachdidaktische Zwecke interessante Korpusrecherchen sind so nicht oder nur mit hohem Aufwand und mit korpuslinguistischer Expertise durchführbar.

Das von der DFG im Programm LIS (Wissenschaftliche Literaturversorgungs- und Informationssysteme) geförderte Projekt ZuMult („Zugänge zu **multi**-modalen Korpora gesprochener Sprache“; Leibniz-Institut für Deutsche Sprache Mannheim / Herder-Institut, Universität Leipzig / HZSK, Universität Hamburg) hat es sich vor diesem Hintergrund zum Ziel gesetzt, für Fremdsprachendidaktiker/-innen entsprechende Such- und Nutzungsmöglichkeiten zu schaffen und zur Verfügung zu stellen.² Die so geschaffenen Suchinstrumente sind flexibel für Recherchen in den dynamisch wachsenden großen Korpora der gesprochenen Sprache nutzbar. So soll einerseits eine gezielte Auswahl von Sprachbeispielen möglich werden, etwa nach schwierigkeitsbezogenen Parametern (u. a. Wortschatzniveau, Standardnähe/-ferne, Sprechgeschwindigkeit, Anteil typisch mündlicher Phänomene). Andererseits werden vielfältige Möglichkeiten für eine didaktische Nutzung der Sprachbeispiele geschaffen (etwa Hervorhebung mündlicher Phänomene im Transkript sowie Bereitstellung ihrer schriftstandardsprachlichen Entsprechungen, etwa *biste – bist du*, Veränderung der Abspielgeschwindigkeit, Download des enthaltenen Wortschatzes, Abgleich mit Vergleichswortschatzlisten u. a.).

² Weitere Informationen zu den Zielen und Aktivitäten des Projekts finden sich auf der Projektwebseite: <https://zumult.org>. Die im Projekt entwickelten Werkzeuge sind sowohl über die DGD als auch über folgenden Link zugänglich: <https://zumult.org/demo/>. Für die Nutzung der im Beitrag vorgestellten Werkzeuge ist die kostenlose Registrierung bei der DGD erforderlich.

Im vorliegenden Beitrag werden anhand eines sprachdidaktischen Anwendungsszenarios Konzeption, Design und Nutzungsmöglichkeiten der neu geschaffenen digitalen Anwendungen vorgestellt.

2 „Ich brauch da dringend ein passendes Beispiel ...“: *ein* Anwendungsszenario

Im Mittelpunkt des hier zugrunde liegenden Anwendungsszenarios steht eine Lehrkraft für Deutsch als Fremdsprache. Sie ist auf der Suche nach einem passenden Beispiel für die Behandlung des Themas „Gespräche beim Backen und Kochen“ im Unterricht und hat dabei bereits eine kurze Recherche auf diversen Internetseiten hinter sich. Zudem war der Versuch, in Korpora ein passendes Beispiel zu finden, erfolglos, da sie sich mit der Nutzung dieser Ressourcen nicht so gut auskennt. Bei der Recherche ist sie jedoch auf ZuMult gestoßen und fragt sich nun, wie es mit Hilfe von den in ZuMult entwickelten Werkzeugen ZuMal und ZuViel möglich ist, ein passendes Gesprächsbeispiel für ihre Lernenden zu finden. In der geplanten Unterrichtseinheit soll es um eine Annäherung an ausgewählte gesprochenessprachliche Merkmale in einem handlungsorientierten Alltagsgespräch gehen. Die Zielgruppe sind Lernende in einem B1-Sprachkurs. Dies sollte bei der Auswahl eines passenden Sprechereignisses vor allem hinsichtlich des Wortschatzes und der Schwierigkeit des Textes für Deutschlernende bedacht werden.³

In den beiden folgenden Abschnitten werden mit ZuMal (Zugang zu Merkmalsauswahl von Transkripten) und ZuViel (Zugang zu Visualisierungs-Elementen für Transkripte) die für dieses Anwendungsszenario relevanten Werkzeuge aus ZuMult jeweils kurz vorgestellt. Es wird zudem exemplarisch gezeigt, wie bei der Suche nach einem Gesprächsbeispiel in ZuMal vorgegangen werden könnte

3 Dem Beitrag liegt ein Videovortrag zugrunde, der im Rahmen der FaDaF-Tagung 2021 an der PH Freiburg in Zusammenarbeit mit der Universität Kassel als Screencast erstellt wurde. In diesem Screencast wird in Form einer kurzen Animation das Anwendungsszenario anhand der fiktiven Lehrkraft Pico eingeführt. Daran anschließend werden Nutzungsmöglichkeiten der im Projekt ZuMult entwickelten Werkzeuge vorgestellt. Der Screencast wird über den FaDaF-Youtube-Kanal online zur Verfügung gestellt: https://www.youtube.com/watch?v=X_wj4Uchb3U&t=3s.

(2.1) und welche Möglichkeiten ZuViel für die Arbeit mit diesem Gespräch bietet (2.2).⁴

2.1 Vorauswahl von Gesprächsbeispielen mit dem Werkzeug ZuMal

Das Werkzeug ZuMal wurde gezielt für die Suche nach passenden Gesprächsbeispielen für didaktische Kontexte entwickelt. Es bietet die Möglichkeit, alle in den beiden Korpora FOLK und GWSS enthaltenen Sprechereignisse nach verschiedenen Merkmalen zu filtern und miteinander im Hinblick auf die Ausprägung der Merkmale zu vergleichen, um etwa die Schwierigkeit verschiedener Sprechereignisse für Lernende genauer einschätzen zu können.

Das Interface von ZuMal (vgl. Abbildung 1) beinhaltet drei Hauptkomponenten, die im Folgenden zunächst vorgestellt werden. Daran anschließend wird bezogen auf die geschilderten Nutzungsinteressen gezeigt, wie ZuMal von Lehrenden bei der Auswahl von Sprechereignissen genutzt werden könnte.

Auf der linken Seite ist die Merkmalsauswahl mit den verschiedenen Auswahlparametern zu sehen (vgl. Abbildung 1, Nr. 1). Diese Merkmalsauswahl ist wiederum in drei Blöcke geteilt, die jeweils unterschiedliche Charakteristika der Sprechereignisse in den Korpora fokussieren. Der erste Block ergibt sich aus den Metadaten zu den Gesprächen. Es ist hier möglich, nach dem Gesprächstyp bzw. nach verschiedenen Interaktionsdomänen (institutionell, privat, öffentlich und Sonstiges) zu filtern. Die einzelnen Interaktionsdomänen sind weiter nach Aktivitäten bzw. Kontexten ausdifferenziert (so etwa institutionell > Bildung > Unterricht oder auch privat > Spielen, privat > Kochen usw.). Außerdem stehen Filter zur Verfügung, die nach der Art des Sprechereignisses (Telefongespräch, Verkaufsgespräch, Tischgespräch, studentisches Alltagsgespräch) und nach den Themen der Sprechereignisse differenzieren. Zwei weitere Filter ermöglichen eine Auswahl anhand einer bestimmten Sprachregion und die Begrenzung der Dauer eines Sprechereignisses.

Der nächste Block an Merkmalen fasst vier schwierigkeitsbezogene Parameter, nämlich Wortschatz, Standardnähe, Sprechgeschwindigkeit und Überlappun-

⁴ Meißner und Wallner (im Erscheinen) stellen in ihrem Beitrag zwei Anwendungsszenarien aus der Perspektive von Lernenden vor, die unabhängig von einem Unterrichtssetting recherchieren und dabei auf die Möglichkeiten von ZuMult zurückgreifen. Dabei wird zunächst ein Szenario geschildert, bei dem es um die Suche nach einem in einer spezifischen Sprachregion verortbaren Sprechereignis geht, und dann ein Szenario zur Suche nach einem Sprechereignis, das bei der Vorbereitung auf eine in Deutschland für ein Studium relevante Gesprächsform hilfreich ist.

gen, zusammen, die sich automatisiert aus den Transkripten der Gespräche berechnen lassen (vgl. hierzu auch Fandrych/Meißner/Wallner; im Druck). Diese beziehen sich dabei auf folgende Konzepte:

Wortschatz: Unter dem Auswahlfilter finden sich verschiedene, den einzelnen Niveaustufen des Gemeinsamen Europäischen Referenzrahmens (GER) zugeordnete Referenzwortlisten. Bei den hier zur Auswahl stehenden Wortlisten handelt es sich um kompetenz- und frequenzbezogene Listen (Goethe-Zertifikatswortschätze für A1, A2 und B1⁵ sowie die auf dem Frequenzwörterbuch von Tschirner/Möhring 2019 beruhenden Frequenzlisten). Es ist so für Lehrende möglich, ein Sprechereignis wortschatzbezogen hinsichtlich des GER-Sprachniveaus einzuschätzen. Über einen Abgleich der Lexik in den Transkripten mit den Referenzwortlisten wird der prozentuale Anteil der Übereinstimmung zwischen Transkript und Wortliste berechnet. Dabei wird davon ausgegangen, dass eine Wortschatzdeckung von 90–95 Prozent die Grundlage für ein ausreichendes Hörverständnis bildet (vgl. Zeeland/Schmitt 2013). Dies ist aber selbstverständlich von den Zielen abhängig, die Lernende mit dem Hören jeweils verbinden (vom Grobverständnis bis zu sehr konkreten und detaillierten Mikrohörübungen; vgl. Dietz 2021).

Standardnähe: Der Auswahlparameter „Standardnähe“ operationalisiert die Normalisierungsrate. Diese zeigt an, wie viele sprachliche Einheiten in einem Gespräch abweichend von einem angenommenen schriftsprachlichen Standard realisiert werden (wie bspw. bei Reduktionen wie *hab* [habe] oder bei Klitisierungen wie *gehste* [gehst du]). Diese Abweichungen können automatisch ermittelt werden, da für alle Gespräche sowohl eine aussprachenah Transkription⁶ als auch eine orthografisch an die Standardlautung angepasste Transkriptversion vorliegen. Ein Normalisierungsfall tritt immer dann auf, wenn die aussprachenah Transkription und die orthografische Standardlautung voneinander abweichen. Eine höhere Normalisierungsrate bzw. ein hoher Wert im Filter „Standardnähe“ deutet daher auf einen hohen Anteil an dialektalen oder auch an umgangssprachlichen Realisierungen hin. Ein niedriger Wert lässt hingegen darauf schließen, dass die sprachlichen Einheiten überwiegend schriftsprachennah realisiert wurden. Allerdings kann anhand der Normalisierungsrate nicht abgeleitet werden, inwieweit es sich bei den Normalisierungsfällen um umgangssprachlich, regionalsprachlich oder zweitsprachlich bedingte Abweichungen handelt.

⁵ Vgl. unter: https://www.goethe.de/pro/relaunch/prf/de/A1_SD1_Wortliste_02.pdf, https://www.goethe.de/pro/relaunch/prf/de/Goethe-Zertifikat_A2_Wortliste.pdf sowie https://www.goethe.de/pro/relaunch/prf/de/Goethe-Zertifikat_B1_Wortliste.pdf (25.11.2021).

⁶ Die Transkripte liegen als aussprachenah Transkription nach dem cGat Minimaltranskript vor (vgl. dazu Schmidt et al. 2015 und den Abschnitt 2.2).

Sprechgeschwindigkeit: Mit diesem Filter kann eine Auswahl bezüglich der pro Sekunde artikulierten Silben in den Gesprächen getroffen werden. Prinzipiell kann davon ausgegangen werden, dass eine höhere Artikulationsrate gerade für Deutschlernende herausfordernder sein kann. Es ist allerdings möglich, über das Werkzeug ZuViel die Sprechgeschwindigkeit beim Abspielen des Gesprächs zu reduzieren oder zu erhöhen (vgl. Abschnitt 2.2). Außerdem sind viele der Sprechereignisse durch einen Wechsel von schneller und langsamer gesprochenen Phasen geprägt.

Überlappungen: Dieser Auswahlfilter ermöglicht es, Gespräche nach dem Anteil überlappend gesprochener Sequenzen einzustufen. Wie bei der Sprechgeschwindigkeit ist beim Anteil gleichzeitig bzw. überlappend gesprochener Sequenzen in den Sprechereignissen davon auszugehen, dass mit einem höheren Anteil zudem höhere Anforderungen an das Hörverstehen der Lernenden einhergehen.

Bei all diesen schwierigkeitsbezogenen Parametern handelt es sich um globale Maße, die das Sprechereignis als Ganzes analysieren. Dies bedeutet, dass etwa in einem Gespräch mit hoher Normalisierungsrate und vielen Überlappungen durchaus Sequenzen vorhanden sein können, die relativ standardnah realisiert wurden und in denen es keine Überlappungen gibt.

Die anderen beiden Filter, die in ZuMal zur Verfügung stehen, sind die Filter „Wortarten“ und „Mündlichkeitsphänomene“. Im ersten dieser Filter kann der Anteil ausgewählter Wortarten in den Sprechereignissen eingestellt werden. Hierzu zählen Nomen (NN), Eigennamen (NE), Verben (V), Adjektive (ADJ) und Adverbien (ADV). Zudem können die Sprechereignisse nach dem Anteil der in Distanzstellung gebrauchten trennbaren Verben (PTKVZ) gefiltert werden. Im Filter „Mündlichkeitsphänomene“ kann eingestellt werden, wie hoch der Anteil an gesprochensprachlichen Merkmalen auf der Textoberfläche sein sollte. Hier stehen Häsitationen (NGHES), Interjektionen, Antwortpartikeln und Interaktionssignale (NGIRR), Modal- und Abtönungspartikeln (PTKMA), Diskursmarker (SEDM), Tag Questions (SEQU) und Klitisierungen (CLITIC) zur Auswahl.

Auf der rechten Seite ist in der oberen Hälfte der ZuMal-Oberfläche ein Streudiagramm zu sehen, das die Ergebnismenge visualisiert (vgl. Abbildung 1, Nr. 2). Auf der x-Achse und auf der y-Achse können dabei jeweils die schwierigkeitsbezogenen Parameter und der Parameter „Dauer“ eingestellt werden. Für den Parameter „Wortschatz“ muss darüber hinaus die Wortliste gewählt werden, die dem Streudiagramm zugrunde gelegt werden soll. Das Streudiagramm passt sich jeweils dynamisch den in den Merkmalsfiltern ausgewählten Einstellungen an, sodass nur die Gespräche im Diagramm auftauchen, die zu den ausgewählten Merkmalen passen. Wird keine Anpassung der Filter vorgenommen, sind alle Gesprächsereignisse innerhalb des Korpus in dem Streudiagramm zu sehen (vgl. Abbildung 1, Nr. 2).

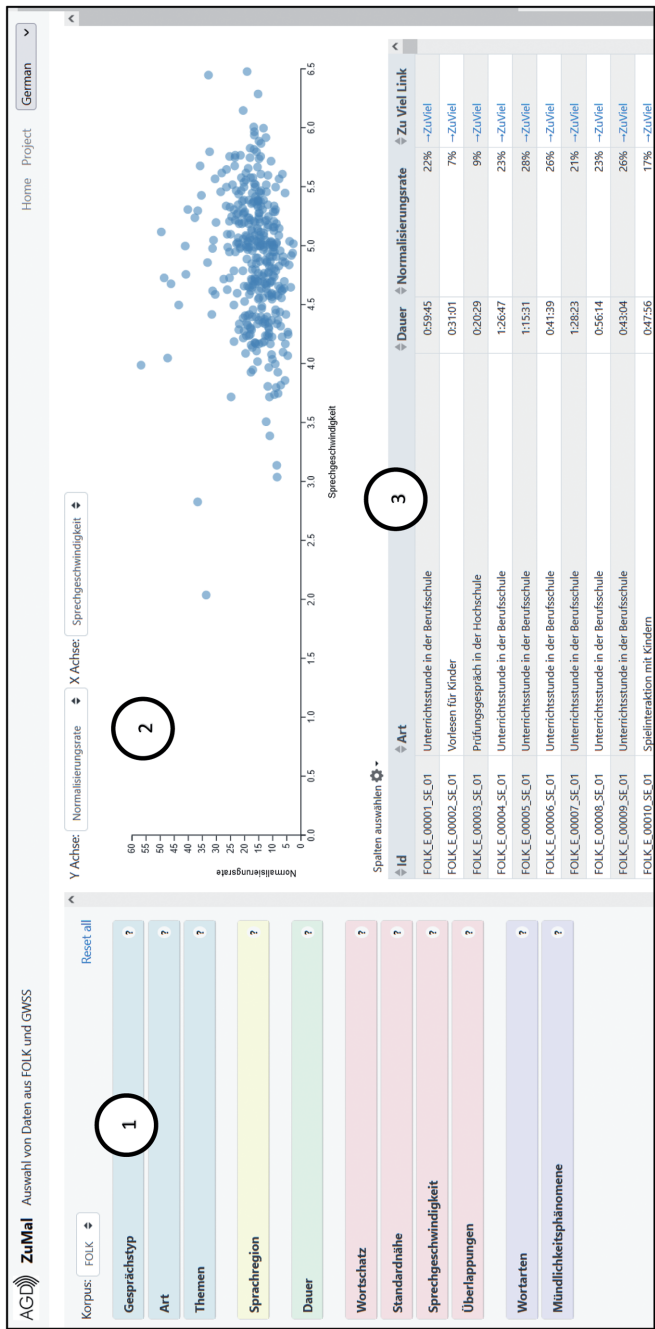


Abbildung 1: ZuMal-Interface

In der unteren Hälfte ist eine bearbeitbare Tabellenansicht zu sehen, die sich ebenfalls dynamisch an die in den Merkmalsfiltern getroffene Auswahl anpasst (vgl. Abbildung 1, Nr. 3). Je nach Merkmalsauswahl wird hier eine Ergebnisliste angezeigt. In der voreingestellten Ansicht sind die ID des Sprechereignisses im ausgewählten Korpus, die Art des Gespräches, die Dauer, die Normalisierungsrate und der Link zum Werkzeug ZuViel, der ein ausgewähltes Transkript in diesem Werkzeug zur weiteren Bearbeitung öffnet (vgl. Abschnitt 2.2), sichtbar. Es können für jeden der verfügbaren Filter die entsprechenden Spalten hinzugefügt oder wieder abgewählt werden, sodass es möglich ist, sich eine den persönlichen Präferenzen angepasste Tabelle zu erstellen.

Für die im Anwendungsszenario dargelegte Suche nach einem passenden Gesprächsbeispiel empfiehlt es sich, das FOLK-Korpus auszuwählen, da es Sprechereignisse aus alltäglichen Kommunikationssituationen bereitstellt. Hier stehen derzeit 374 Gespräche zur Verfügung.⁷ Im nächsten Schritt wird der Gesprächstyp eingeschränkt. Hier wird als Interaktionsdomäne „privat“ ausgewählt. Dies reduziert die Auswahl an Gesprächen auf 140. Daraufhin bietet es sich an, mit Hilfe der Gruppierung nach Aktivitäten unterhalb der Auswahlenebene „privat“ eine weitere Einschränkung vorzunehmen. Dem eingangs formulierten Wunsch der Lehrperson, dass es sich um ein Gespräch beim Kochen oder Backen handeln sollte, entsprechen die folgenden Sprechereignisse: „Kommunikation beim Kochen“ (vier Gespräche), „Backinteraktion“ (ein Gespräch), „Kochinteraktion“ (ein Gespräch) und „Backen mit Freunden“ (ein Gespräch). Die Filter „Themen“ und „Sprachregion“ werden in diesem Anwendungsszenario nicht benötigt. Der nächste Filter, der bei der weiteren Auswahl eine Rolle spielt, ist der Filter „Dauer“. Die sieben Gespräche zeichnen sich durch eine relativ unterschiedliche Länge aus, wobei das kürzeste Sprechereignis eine Dauer von 18:56 Minuten und das längste eine Dauer von 1:26 Stunden hat. Drei der sieben Gespräche dauern über eine Stunde und sind daher aufgrund ihrer Länge für die Unterrichtsvorbereitung möglicherweise nur bedingt geeignet.⁸ Mit Hilfe des Filters „Dauer“ werden daher alle Gespräche ausgeschlossen, die länger als eine Stunde dauern. Die verbleibenden vier Gespräche („Kommunikation beim Kochen“ [FOLK_E_00300_SE_01 und FOLK_E_00225_SE_01], „Backen mit Freunden“ [FOLK_E_00372_SE_01] und „Kochinteraktion“ [FOLK_E_00332_SE_01]) sollen nun anhand der schwierigkeitsbezogenen Filter auf ihre Eignung hinsichtlich eines Einsatzes im Unterricht überprüft werden (vgl. hierzu die Ergebnistabelle in Abbildung 2, in der die vier verblei-

7 Stand: 13.7.2021, Version 2.16: Release am 17.5.2021.

8 Es ist allerdings auch möglich, bei der Arbeit mit dem Werkzeug ZuViel einzelne Sequenzen aus den Sprechereignissen auszuwählen (vgl. Abschnitt 2.2).

benden Sprechereignisse zusammen mit den schwierigkeitsbezogenen Parametern abgebildet sind). Das vorliegende Anwendungsszenario bezieht sich auf eine Lernendengruppe in einem B1-Sprachkurs Deutsch. Deshalb wird über den Filter „Wortschatz“ zunächst die Referenzwortliste „GOETHE_B1“ ausgewählt. Wie bereits angesprochen, kann davon ausgegangen werden, dass für ein ausreichendes Hörverstehen eine Wortschatzdeckung von mindestens 90 Prozent gegeben sein sollte. Alle vier Gespräche weisen eine relativ hohe Wortschatzdeckung um die 90 Prozent auf. Lediglich das Sprechereignis „Kochinteraktion“ bleibt mit 89 Prozent Wortschatzdeckung knapp unter dem geforderten Schwellenniveau (vgl. Abbildung 2 „Goethe B1 Deckung“) und wird daher für die folgenden Schritte ausgeklammert.

Spalten auswählen ⚙

Index	Id	Art	Dauer	Normalisierungsrate	Überlappungen	Artikulationsrate	Goethe B1 Deckung	Zu Viel Link
1	FOLK_E_00332_SE_01	Kochinteraktion	0:18:56	19%	1.69	5.03	89%	→ZuViel
2	FOLK_E_00372_SE_01	Backen mit Freunden	0:29:20	36%	8.93	5.30	90%	→ZuViel
3	FOLK_E_00300_SE_01	Kommunikation beim Kochen	0:31:54	17%	48.03	4.71	91%	→ZuViel
4	FOLK_E_00225_SE_01	Kommunikation beim Kochen	0:48:16	17%	11.58	5.05	92%	→ZuViel

Abbildung 2: Ergebnisliste in ZuMal mit Vergleich von schwierigkeitsbezogenen Parametern

Die drei übrigen Transkripte weisen alle eine ähnlich hohe Artikulationsrate von etwa fünf artikulierten Silben pro Sekunde auf. Allerdings unterscheiden sie sich deutlich im Hinblick auf die Normalisierungsrate und den Wert für Überlappungen. So variiert der Wert für Überlappungen zwischen 8,93 bei dem Gespräch „Backen mit Freunden“ (FOLK_E_00372_SE_01), 11,58 für „Kommunikation beim Kochen“ (FOLK_E_00225_SE_01) und dem relativ hohen Wert von 48,03 bei dem Gespräch „Kommunikation beim Kochen“ (FOLK_E_00300_SE_01). Hinsichtlich des Anwendungsszenarios scheint dieses Sprechereignis aufgrund der hohen Überlappungsrate für Lernende auf dem Niveau B1 nur bedingt geeignet, da davon auszugehen ist, dass ein höherer Anteil an überlappend gesprochenen Äußerungssequenzen das Verständnis deutlich erschwert. Gleichzeitig weist „Backen mit Freunden“ im Vergleich zu den anderen Gesprächen eine deutlich höhere Normalisierungsrate von 36 Prozent auf. Dies lässt vermuten, dass es sich hier um ein stärker umgangssprachlich, möglicherweise auch regionalsprachlich geprägtes Sprechereignis handelt.⁹

⁹ Abschnitt 2.2 enthält ausgewählte Belegstellen aus diesem Sprechereignis, an denen zahlreiche Normalisierungsfälle deutlich werden, vgl. Beleg (1) sowie die Belege (3) bis (8).

Da die Lehrkraft im eingangs beschriebenen Anwendungsszenario in ihrem Unterricht genau solche Phänomene thematisieren möchte, erscheint dieses Gespräch besonders geeignet.

Über den ZuViel-Link in der Ergebnistabelle kann nun das dazugehörige Transkript in ZuViel geöffnet und weiterbearbeitet werden.

2.2 Vorbereitung und Nutzung von Gesprächsbeispielen im Unterricht mit dem Werkzeug ZuViel

Mit ZuViel (Zugang zu Visualisierungs-Elementen für Transkripte) wurde ein Werkzeug entwickelt, das vielfältige didaktische Arbeitsmöglichkeiten mit einem zuvor über die Merkmalsauswahl in ZuMal ausgewählten Sprechereignis ermöglicht. Das Interface umfasst vier Hauptkomponenten, welche hier im Folgenden am Beispiel des ausgewählten Sprechereignisses „Backen mit Freunden“ (FOLK_E_00372_SE_01) erläutert werden.

Auf der rechten Seite (vgl. Abbildung 3, Nr. 1) befinden sich zwei Videomitschnitte des Sprechereignisses, die aus unterschiedlichen Perspektiven aufgenommen wurden. Beim Abspielen wird im oberen Video die Transkription mit angezeigt. Bei Gesprächen, in denen kein Video zur Verfügung steht, findet sich an dieser Stelle die Audioaufnahme. Über den Play-Button kann das Abspielen des Videos bzw. der Audioaufnahme gestartet werden. Falls die Sprechgeschwindigkeit für die Lernenden zu hoch sein sollte, besteht die Möglichkeit, über den Menüpunkt „Parameter“ die Abspielgeschwindigkeit zu reduzieren, ohne die Stimmhöhe zu verändern.

In der Mitte (vgl. Abbildung 3, Nr. 2) wird das Transkript des Sprechereignisses angezeigt. Voreingestellt ist die aussprachenähe Transkription. Diese erfolgte nach den Konventionen für computergestützte GAT-Minimaltranskripte in literarischer Umschrift (Schmidt/Schütte/Winterscheidt 2015). Neben einer generellen Kleinschreibung und dem Verzicht auf Interpunktion bedeutet das, dass alle Abweichungen von der Standardlautung notiert werden. Dazu zählen etwa Reduktionen (z. B. *jetz* [jetzt] in Beitrag 0001, *hab* [habe] in Beitrag 0006 sowie *un* [und] in Beitrag 0017), Klitisierungen (z. B. *hamma* [haben wir] in Beitrag 0388 und *geht s* [geht es] in Beitrag 0434), regionale Einflüsse (z. B. *budder* [Butter] in Beitrag 0023, *uffgemacht* [aufgemacht] in Beitrag 0030) oder auch typisch mündliche Schnellsprechformen (z. B. *ham* [haben] in Beitrag 0806). Auch Hässitationen (*äh*), Pausen (z. B. *(.)* für eine Mikropause von bis zu 0,2 Sekunden Dauer), Atmen (*°h* für Einatmen, *h°* für Ausatmen) und parasprachliche Merkmale wie Lachen oder Schmatzen wurden bei der Transkription festgehalten. Dies ermöglicht eine relativ gute Annäherung an die gesprochene Sprache, wie bereits der erste Beitrag aus dem Sprechereignis „Backen mit Freunden“ zeigt:

- (1) okay (.) dann fange ma jetzt mo an (FOLK_E_00372_SE_01, Beitrag 0001)

Neben der aussprachenahen Transkription kann über den Menüpunkt „Parameter“ auch eine orthografisch normalisierte Fassung des Transkripts aufgerufen werden. Auch alle abweichend realisierten sprachlichen Einheiten werden dann in dieser Ansicht in ihrer standardorthografischen Fassung angezeigt. Der folgende Beleg zeigt dies zur Illustration für den in (1) aufgeführten aussprachenahen Beleg:

- (2) okay (.) dann fangen wir jetzt mal an (FOLK_E_00372_SE_01, Beitrag 0001)

Bei Bedarf kann auch eingestellt werden, dass diejenigen sprachlichen Einheiten, die abweichend von der standardorthografischen Lautung realisiert wurden, im Transkript durch Unterstreichung hervorgehoben werden. So kann die Aufmerksamkeit der Lernenden zusätzlich auf diese Einheiten gelenkt werden, was insbesondere für die Entwicklung und Förderung von Sprachbewusstheit und rezeptiver Varietätenkompetenz sehr sinnvoll sein kann. Denkbar ist hier einerseits, dass den Lernenden die aussprachenaher Transkription an die Hand gegeben wird mit dem Auftrag, die standardorthografische Form der unterstrichenen Einheiten selbstständig (oder auch unter Zuhilfenahme einer bereitgestellten Liste) zu ergänzen. Andererseits könnte auch mit der standardorthografischen Fassung gearbeitet werden und die Lernenden ergänzen die aussprachenahere Entsprechung zu den unterstrichenen Einheiten. Die Bereitstellung der beiden Fassungen ermöglicht zudem auch binnendifferenziertes Arbeiten. Während leistungstärkere Lernende mit der aussprachenahen Fassung arbeiten, erhalten die leistungsschwächeren Lernenden entweder ergänzend oder ausschließlich die orthografisch normalisierte Version des Transkripts.

Zusätzlich ist es auch möglich, sich anstelle der aussprachenahen Transkription die Wortart (bspw. *wir* = Personalpronomen [PPER], *mal* = Modalpartikel [PTKMA]), die Grundformen (*uffgemacht* – *aufmachen*) oder auch die phonologische Umschrift (ʔo.ke: (.) dan faŋ.ə ma jɛts mo ʔan) anzeigen zu lassen.

Darüber hinaus ist über das Suchfenster auf der rechten Seite in der Menüleiste ein gezieltes Ansteuern sprachlicher Einheiten im Transkript möglich. Möchten Lehrende beispielsweise die Stellen im Transkript finden, in denen Modalpartikeln vorkommen, dann kann in das Suchfenster „pos=PTKMA“ eingegeben werden. Unabhängig davon, welche Transkriptansicht eingestellt ist (aussprachenaher, orthografisch normalisierte usw.), werden im Transkript alle Modalpartikeln markiert. Im ausgewählten Sprechereignis „Backen mit Freunden“ sind dies unter anderem *mal* (auch realisiert als *ma*, *mo* oder *mol*, vgl.

auch Beleg [1]), *denn* (auch realisiert als *n*, vgl. Beleg [3]), *aber* (vgl. Beleg [4]), *schon* (auch realisiert als *scho* vgl. Beleg [5]) und *ja* (vgl. Beleg [6]).

- (3) [wo is **n** der schneebesen] (FOLK_E_00372_SE_01_T_01, Beitrag 0143)
- (4) isch glab der muss **aber** nit so lang in backofen (FOLK_E_00372_SE_01_T_01, Beitrag 0323)
- (5) des pass **scho** (FOLK_E_00372_SE_01_T_01, Beitrag 0572)
- (6) kannscht **ja** morje kuche mit auf die arbeit nehme lisa (FOLK_E_00372_SE_01_T_01, Beitrag 0680)

Neben Modalpartikeln können selbstverständlich auch andere Wortarten angesteuert werden. Eine Übersicht über die Kürzel, die die einzelnen Wortarten repräsentieren, lässt sich über das Icon links neben dem Suchfenster aufrufen. Ebenso kann im Transkript nach bestimmten orthografisch normalisierten Formen gesucht werden. Interessieren sich Lehrende etwa für alle aussprachenah (und nicht standardorthografisch) transkribierten Varianten von *aber*, können sie in das Suchfenster „norm=aber“ eingeben. In der aussprachenahen Transkriptansicht finden sich dann Belege, in denen *aber* als *aba* (vgl. Beleg [7]) oder auch als *awa* (vgl. Beleg [8]) realisiert wird:

- (7) [**aba** die]müsse ma noch a bissche ufftaue lasse weil die so hart war (FOLK_E_00372_SE_01_T_01, Beitrag 0037)
- (8) [dann]müsse ma **awa** [die]äh backforme] vom owe runna hole (FOLK_E_00372_SE_01_T_01, Beitrag 0151)

Durch Klicken in eine beliebige Stelle des Transkripts wird die Video- bzw. Audioaufnahme ab der ausgewählten Position abgespielt.

Auf der linken Seite (vgl. Abbildung 3, Nr. 3) wird eine automatisch generierte Liste bereitgestellt, die den im Sprechereignis vorkommenden Wortschatz mit Angaben zur Frequenz umfasst. Ergänzend dazu kann eine Referenzwortliste aufgerufen werden, die anzeigt, ob der in der Referenzwortliste enthaltene Wortschatz auch im Sprechereignis vorkommt. Ist dies der Fall, erhalten die entsprechenden sprachlichen Einheiten ein Häkchen. Zusätzlich wird dieser Wortschatz auch im Transkript gelb unterlegt. Bei den hier zur Auswahl stehenden Referenzwortlisten handelt es sich einerseits um die auch in ZuMal für die Vorauswahl der Sprechereignisse genutzten kompetenz- und frequenzbezogenen Listen (vgl. 2.1). Lehrende können so beispielsweise überprüfen, welcher Wortschatz des Sprechereignisses ihren Lernenden auf einer bestimmten Niveaustufe voraussichtlich bekannt ist und welcher Wortschatz gegebenenfalls vorentlastet werden müsste. Andererseits stehen thematische Listen zur Auswahl. Diese umfassen die Lexik der Wortschatzbereiche „Essen“, „Haus und Wohnung“ und „Schule und Ausbildung“ und basieren auf dem Übungswortschatz „Sage und Schreibe“ (Fandrych/

Tallowitz 2019). Lehrende können so überprüfen, welcher themenbezogene Wortschatz im Sprechereignis enthalten ist. Durch Klicken auf die Einheiten in der Liste können die entsprechenden Stellen im Transkript direkt angesteuert werden. So sind beispielsweise im Sprechereignis mehrere Maßeinheiten wie *Gramm* und *Milliliter* enthalten. Möchten Lehrende diese im Kontext thematisieren, ließen sich die entsprechenden Stellen direkt aufrufen. Klickt man etwa auf *Gramm*, gelangt man direkt an eine Stelle, in der gerade etwas abgemessen wird (vgl. Abbildung 4).

Für die Thematisierung im Unterricht ist es möglich, eine beliebige Sequenz über das graue Icon am Zeilenbeginn (via Start bzw. *End selection*) auszuwählen, in einem extra Tab gesondert aufzurufen und bei Bedarf herunterzuladen. Der Download umfasst dabei sowohl den ausgewählten Transkriptausschnitt in verschiedenen Formaten¹⁰ als auch die auf diesen Ausschnitt begrenzte Video- bzw. Audioaufnahme. Natürlich ist es auch möglich, das gesamte Transkript zuzüglich der Video- bzw. Audioaufnahme herunterzuladen. Ergänzend dazu stehen verschiedene Wortlisten zum Download zur Verfügung. Zur Auswahl stehen hier Wortlisten, die

- a) die gesamte im Transkript enthaltene Lexik enthalten,
- b) den durch eine ausgewählte Referenzwortliste gedeckten oder wahlweise auch nicht gedeckten Wortschatz umfassen,
- c) alle vom schriftsprachlichen Standard abweichend realisierten Einheiten in ihrer aussprachenahen Transkription zzgl. ihrer normalisierten und lemmatisierten Fassungen enthalten.

Eine weitere Komponente, der *Density Navigator*, befindet sich im Interface auf der rechten Seite oberhalb der Video- bzw. Audioaufnahme (vgl. Abbildung 2, Nr. 4). Durch Klicken auf das Pluszeichen in der rechten oberen Ecke kann er vergrößert werden. Der *Density Navigator* gibt eine kompakte Übersicht über die zeitliche Anordnung der Gesprächsbeiträge der verschiedenen Sprecher/-innen. Für jede Sprecherin bzw. jeden Sprecher steht ein farbiger Balken. Bei dem Sprechereignis „Backen mit Freunden“ gibt es beispielsweise drei verschiedene Sprecher/-innen und entsprechend drei Balken. Anhand des *Density Navigators* ist erkennbar, dass zu Beginn des Sprechereignisses zunächst nur zwei der drei Sprecher/-innen aktiv sind. Der dritte Sprecher setzt erst später mit ein. Durch diese Art der Visualisierung des Gesprächsverlaufs lassen sich Passagen mit verstärkter Interaktivität oder auch überlappende Äußerungssequenzen auf einen Blick leichter

¹⁰ Darunter auch txt, welches für die Weiterverarbeitung in Word zu empfehlen ist.

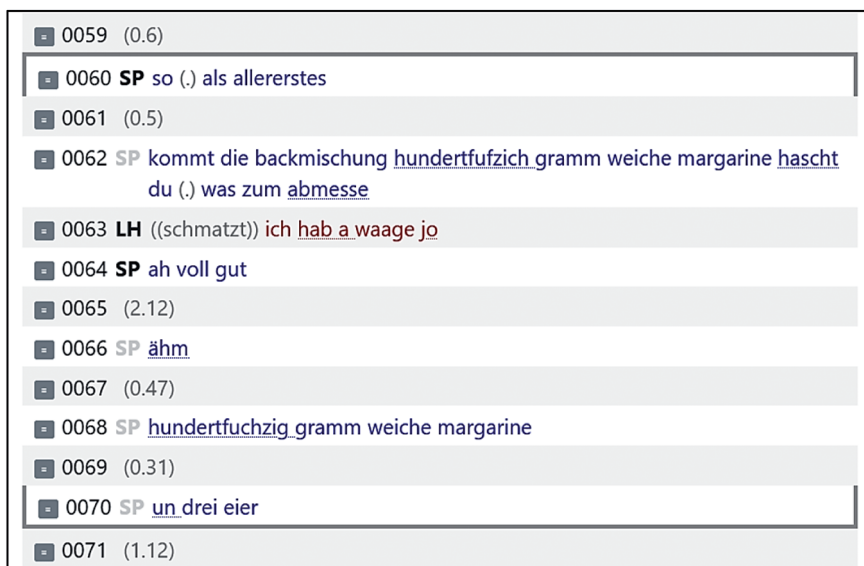


Abbildung 4: Ausgewählte Sequenz mit Kontext zum Suchwort *Gramm* aus FOLK_E_00372_SE_01

identifizieren und von eher monologischen Passagen unterscheiden. Anhand der Farbintensität der einzelnen Balken ist im *Density Navigator* zudem erkennbar, inwieweit die einzelnen sprachlichen Einheiten von der orthografischen Standardlautung abweichend realisiert wurden. Je heller die Farbe, desto stärker sind die Abweichungen. Zudem wird innerhalb der Balken mit Hilfe der Anzahl der Pfeile die Sprechgeschwindigkeit visualisiert. Drei Pfeile (>>>) stehen für eine hohe Sprechgeschwindigkeit, zwei Pfeile (>>) kennzeichnen eine mittlere und ein Pfeil (>) eine niedrige Sprechgeschwindigkeit. Lehrende haben so die Möglichkeit, sich einen Überblick über das Sprechereignis zu verschaffen und je nach Bedarf standardnahe oder abweichend realisierte bzw. schnellere oder langsamere Äußerungspassagen für ihre Lernenden zu identifizieren.

3 Fazit und Ausblick

Wie eingangs erwähnt, stellen die schnell wachsenden Bestände gesprochen-sprachlicher Daten und Sprechereignisse ein großes Reservoir für die Sprachdidaktik dar, das bisher noch kaum genutzt worden ist. Dies liegt auch daran, dass die Nutzungsforschung für solche – mit hohem Aufwand erarbeiteten – Ressourcen noch am Anfang steht. In einer empirischen Vorstudie zu ZuMult (vgl. Fandrych et al. 2016) wurde unter den bestehenden Nutzer/-innen der DGD, des Kor-

pus GWSS und der vom HZSK angebotenen Korpora ermittelt, dass die meisten registrierten Nutzer/-innen keine ausgewiesenen Korpuslinguist/-innen sind und viele der eher technischen Suchfunktionen nicht kannten oder mit ihnen nicht hinreichend vertraut waren. Gleichzeitig war die Mehrheit der Nutzer/-innen jedoch in verschiedenen sprachdidaktischen Anwendungskontexten tätig (und nicht in der sprachwissenschaftlichen Forschung oder im Bereich der Korpuslinguistik). Als ein wichtiger Anwendungsbereich erwies sich die Sprachvermittlung. Ziel von ZuMult ist es, diese Kluft zu überbrücken, sodass die vorhandenen Ressourcen möglichst umfassend, aber auch intuitiv und zielgerichtet für sprachdidaktische Zwecke zugänglich werden. Wie bereits angesprochen, sind die hier vorgestellten Tools und Arbeitsmöglichkeiten generischer Art, das heißt, sie lassen sich auch auf in der Zukunft noch in die DGD neu aufgenommene Korpora oder Korpuserweiterungen anwenden. So kann mit ZuMult zumindest exemplarisch aufgezeigt werden, wie eine nutzerorientierte Korpusaufbereitung und entsprechende Tools aussehen könnten, auch wenn es hier noch eine ganze Reihe von Desiderata gibt. Diese beziehen sich zum einen auf grundlegende linguistische und lernbezogene Konstrukte und Konzepte, die als Basis für Operationalisierungen dienen können. Hierzu gehört etwa das Konstrukt „Schwierigkeit“, das neben morphosyntaktischen, phonetischen, semantisch-lexikalischen und interaktionalen auch pragmatisch-situative Faktoren beinhalten müsste und auch nicht lerner- und lernsituationsunabhängig zu bestimmen ist. Letztlich ist damit das Konstrukt nicht wirklich allgemein für alle Kontexte und Lerner/-innen operationalisierbar. Für eine möglichst deutliche Annäherung bedarf es aber eben nicht nur der Einbeziehung weiterer linguistischer Parameter, sondern auch empirischer Forschung zu lernerseitigen Verstehensschwierigkeiten. Daneben stellt die weitere Ausarbeitung von Typologien von Gesprächsereignissen und Gesprächsarten sowie von pragmatischen Kategorien auf der Mesoebene (etwa Erklärsequenzen, argumentative Sequenzen, Scherz- und Konfliktsequenzen) ein wichtiges Desiderat dar; auch ein umfassender Begriff davon, was unter „gesprochenem Standarddeutsch“ zu verstehen ist, würde helfen, die hier als Hilfskonstruktion verwendete Normalisierungsrate zu ergänzen oder zu ersetzen. Daneben war es aufgrund der Förderrichtlinien der LIS-Projekte der DFG (es durften keine weiteren Annotationen oder Datenaufbereitungen durchgeführt werden) und der begrenzten Laufzeit des Projekts (es endet im März 2022) auch im Rahmen von ZuMult notwendig, sich im Projekt auf Aufbereitungswege und Tools zu beschränken, welche die in den Daten bereits angelegten Informationen auf möglichst sinnvolle und nutzerorientierte Weise auffindbar machen und anwendungsorientierte Arbeitsmöglichkeiten bieten. So konnten etwa keine weiteren Annotationen durchgeführt werden (etwa nach bestimmten Gesprächsphasen wie Eröffnungen, Kernphasen und Gesprächsbeendigungen oder nach Themensträngen und Unter-

themen). Da viele Gespräche eine beträchtliche Länge aufweisen, sind diese in vielerlei Hinsicht in sich heterogen. Um gezielter nach relevanten Passagen suchen zu können, wären weitere Aufbereitungsschritte notwendig, die im Rahmen des Projekts nicht durchgeführt werden konnten.

Es bleibt zu hoffen, dass die in ZuMult entwickelten Tools und Anwendungsmöglichkeiten dennoch von der Fachcommunity intensiv genutzt und erprobt werden und diese dann in Folgeprojekten weiterentwickelt und ergänzt werden können.

Literatur

- Dietz, Gunther (2021): „Fremdsprachliches Hörverstehen: Schwächen der traditionellen Hörverstehensdidaktik – Perspektiven der Vermittlung für Deutsch als Fremdsprache“. In: *Deutsch als Fremdsprache* 58 (2), 67–75.
- Fandrych, Christian; Frick, Elena; Hedeland, Hanna; Iliash, Anna; Jettka, Daniel; Meißner, Cordula; Schmidt, Thomas; Wallner, Franziska; Weigert, Kathrin; Westpfahl, Swantje (2016): „User, who art thou? User Profiling for Oral Corpus Platforms“. In: Calzolari, Nicoletta; Choukri, Khalid; Declerck, Thierry; Goggi, Sara; Grobelnik, Marko (Hrsg.): *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC 2016), Portorož, Slovenia*. Paris: European Language Resources Association (ELRA), 280–287.
- Fandrych, Christian; Tallowitz, Ulrike (2019): *Sage und Schreibe: Übungswortschatz Grundstufe A1–B1 mit Lösungen*. Neubearbeitung mit Audio-CD, 1. Auflage. Stuttgart: Ernst Klett.
- Fandrych, Christian; Meißner, Cordula; Wallner, Franziska (im Druck): „Korpora gesprochener Sprache und Deutsch als Fremd- und Zweitsprache: eine chancenreiche Beziehung“. *Korpora Deutsch als Fremdsprache* 1 (2).
- Günthner, Susanne; Wegner, Lars; Weidner, Beate (2013): „Gesprochene Sprache im DaF-Unterricht – Möglichkeit der Vernetzung der Gesprochene-Sprache-Forschung mit der Fremdsprachenvermittlung“. In: Moraldo, Sandro; Missaglia, Federica (Hrsg.): *Gesprochene Sprache im DaF-Unterricht: Grundlagen – Ansätze – Praxis*. Heidelberg: Winter, 113–150.
- Liedke, Martina (2013): „Mit Transkripten Deutsch lernen“. In: Moraldo, Sandro; Missaglia, Federica (Hrsg.): *Gesprochene Sprache im DaF-Unterricht: Grundlagen – Ansätze – Praxis*. Heidelberg: Winter, 243–266.
- Meißner, Cordula; Wallner, Franziska (im Erscheinen): „Korpora gesprochener Sprache als virtuelle Lernräume der Mündlichkeitsdidaktik: Affordanzen eines außerunterrichtlichen Sprachlernsettings“. Erscheint in: Feick, Diana; Rymarczyk, Jutta (Hrsg.): *Fremdsprachenunterricht im virtuellen Raum*. Bern: Peter Lang.
- Merklin, Martina; Riedel, Andrea (2016): „Gesprochenes Deutsch im DaF-Unterricht: Authentizität versus Schriftstruktur“. In: Handwerker, Brigitte; Bäuerle, Rainer; Sieberg, Bernd (Hrsg.): *Gesprochene Fremdsprache Deutsch*. Baltmannsweiler: Schneider, 201–217.
- Schmidt, Thomas; Schütte, Winfried; Winterscheidt, Jenny (2015): *cGAT: Konventionen für das computergestützte Transkribieren in Anlehnung an das Gesprächsanalytische Transkriptionssystem 2 (GAT2)*. Online: <https://ids-pub.bsz-bw.de/frontdoor/index/index/docId/4616> (27.07.2021).

- Tschirner, Erwin; Möhring, Jupp (2019): *A Frequency Dictionary of German: Core vocabulary for learners*. 2. Auflage. London: Routledge.
- Zeeland, Hilde van; Schmitt, Norbert (2013): „Lexical Coverage in L1 and L2 Listening Comprehension: The Same or Different from Reading Comprehension?“. In: *Applied Linguistics*. 457–479. Online: DOI: /10.1093/applin/ams074 (05.10.2021).

Biographische Angaben

Christian Fandrych

ist Professor für Deutsch als Fremdsprache mit Schwerpunkt Linguistik am Herder-Institut, Universität Leipzig. Er arbeitet unter anderem zu den Themen Lexik- und Grammatikvermittlung, Korpuslinguistik, Text- und Gesprächslinguistik, Wissenschaftssprache und Sprachenpolitik. Er ist Chefredakteur der Zeitschrift *Deutsch als Fremdsprache* und Mitherausgeber der Zeitschrift *Fremdsprache Deutsch*.

Matthias Schwendemann

ist wissenschaftlicher Mitarbeiter am Herder-Institut der Universität Leipzig. Seine Arbeitsschwerpunkte sind korpuslinguistische Zugänge zur Erforschung der Entwicklung und des Erwerbs des Deutschen als Fremd- und Zweitsprache und Deutsch als fremde Wissenschaftssprache. Er ist Mitglied der Redaktion der Zeitschrift *Deutsch als Fremdsprache*.

Franziska Wallner

ist wissenschaftliche Mitarbeiterin am Herder-Institut der Universität Leipzig. Ihre Forschungsschwerpunkte sind unter anderen das Deutsche als fremde Bildungs- und Wissenschaftssprache, die korpusbasierte Erforschung der gesprochenen Sprache, Mündlichkeitsdidaktik sowie die Nutzung von Korpora im Kontext von Deutsch als Fremd- und Zweitsprache. Sie ist Mitglied der Redaktion der Zeitschrift *Deutsch als Fremdsprache*.