

5. Operationalisierung: Hegemonial strukturierte Fixierungen von Differenzbeziehungen in sprachlichen Artikulationen

Kapitel 5 befasst sich mit der ausführlichen Darstellung, der in der Fallstudie angewandten Methoden zur Erforschung der diskursiven Formation Globaler Bildung in den Global Governance-Strukturen, im Anschluss an die im Kapitel 4 dargestellte Hegemonie- und Diskurstheorie. Laclau und Mouffe haben sich selbst kaum zur empirischen Umsetzung ihrer Hegemonie- und Diskurstheorie geäußert (Glasze 2013b: 97). Empirische Studien, die dennoch auf ihrer Theorie aufbauen wollen, finden eine mögliche Form der Operationalisierung bei dem aus der Disziplin der Humangeografie kommenden Georg Glasze (Glasze 2007, 2008, 2013b). Dieser bietet eine »Übersetzung« ihrer Diskurstheorie in die Begrifflichkeiten ausgewählter sprach- und literaturwissenschaftlicher Verfahren an: der Lexikometrie und der Analyse narrativer Muster. Beide Verfahren untersuchen die temporäre Fixierung von Bedeutung in Texten (Glasze 2013b: 98). Er begründet diese Möglichkeit der Übersetzung damit, dass »sowohl die Diskurstheorie als auch die Analyseverfahren vor dem Hintergrund strukturalistischer Ansätze und deren Radikalisierung im Poststrukturalismus entwickelt wurden« (ebd.: 98). Die Darstellung der Form der Operationalisierung wird dabei in einer Weise artikuliert, die dem empirischen Material der Fallstudie und ihren Zielen angepasst ist.

Von den wenigen empirischen Beiträgen in der Erziehungswissenschaft, die im Anschluss an die Theorie Laclaus und Mouffes entwickelt wurden, wird wenig umfangreich auf die methodische Umsetzung eingegangen (Macgilchrist 2015; Jergus 2015; Chronaki/Kollosche 2019). Dies mag einerseits darauf zurückzuführen sein, dass Laclaus und Mouffes politische Diskurstheorie bisher zum größten Teil in Analysen in eher soziologischen und politikwissenschaftlichen Feldern zur Anwendung kam und deshalb kein großes Interesse an diesem theoretischen Format der Analyse zu bestehen schien und eine umfangreiche Beschäftigung mit ihnen als empirischer Werkzeugkasten in der Erziehungswissenschaft erst noch aussteht. Andererseits kann dies aber auch mit der Schwierigkeit im Zusammenhang stehen, wie die Laclau-Schüler Howarth und Stavrakakis betonen, dass in empirischen Fallstudien im Anschluss an die politische Diskurstheorie nicht einfach die Konzepte und Logiken des theoretischen Rahmens dem empirischen

Material aufgezungen werden können, sondern sich die Theorie in einer Art artikulieren sollte, die dem jeweiligen empirischen Material fallspezifisch gerecht wird. Das setzt in jedem Fall der Anwendung der Diskurstheorie eine gewisse Offenheit und Flexibilität gegenüber dem Prozess der Anwendung der Theorie am Material voraus (Howarth/Stavrakakis 2000: 4f). Aufgrund des hiesigen Werkzeugkastens, den die Diskurstheorie bereitstellt, des hohen Abstrahierungsgrades ihrer ontologischen Kategorien, der erforderten forschenden Sensibilität für das Changieren zwischem Ontischen und Ontologischen sowie der kaum noch überschaubaren methodischen Möglichkeiten in der erziehungswissenschaftlichen Diskursforschung (bspw. Fegter et al. 2015: 29), die zu ihrer Operationalisierung dienen könnten, scheint dies nicht verwunderlich.

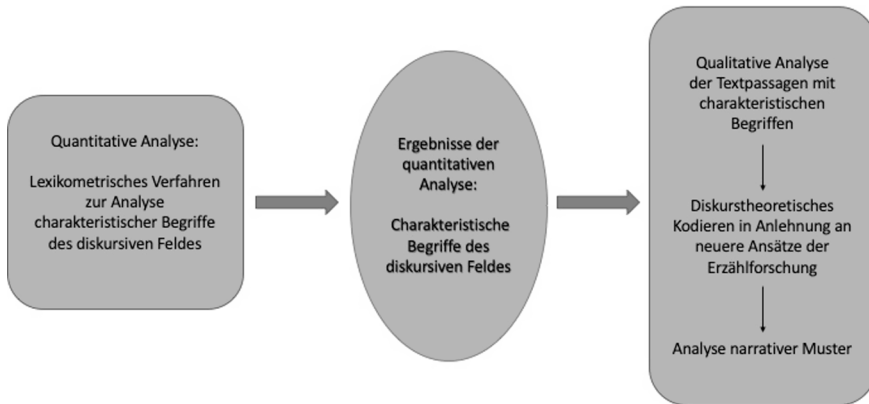
5.1 Triangulation zweier etablierter Sets an Verfahren

In Anlehnung an die Operationalisierung der Diskurstheorie nach Laclau und Mouffe wie sie Glasze vorschlägt (Glasze 2007, 2008, 2013b), haben wir es im Zuge der Umsetzung mit einer Triangulation¹ zweier etablierter Sets an Verfahren zu tun: Der Lexikometrie (Quantitativ) und der Analyse narrativer Muster (Qualitativ) (Glasze 2008: 195). Quantitative und qualitative Analysearten als Verfahrensschritte in einem übergeordneten Forschungsdesign miteinander zu kombinieren, stellt nach Mayring einen hohen Grad an Integration beider Verfahren ineinander dar. Wie er darlegt, gibt es dazu auf der Designebene unterschiedliche Möglichkeiten, quantitative und qualitative Analysearten miteinander zu verbinden. Die in dieser Arbeit veranschaulichte operationalisierte hegemonietheoretische Diskursanalyse, stellt auf der Designebene nach Mayring ein Vertiefungsmodell dar (Mayring 2001: 8). Hierbei wird das als erster Untersuchungsschritt abgeschlossene quantitative lexikometrische Verfahren durch den qualitativen Untersuchungsschritt der Analyse narrativer Muster weitergeführt (siehe auch Abb. 4). Den Ergebnissen der lexikometrischen Analyse und ihrer Herausstellung charakteristischer Be-

1 Unter Triangulation wird nach Norman Denzin herkömmlicherweise der Gebrauch mehrerer Methoden zur Untersuchung des selben Gegenstands verstanden (Denzin 1978: 294). Triangulationen können uns helfen ein tieferes Verständnis von dem befragten Phänomen zu erlangen (Denzin 2012: 82). Laut Denzin gibt es wiederum vier grundlegende Typen der Triangulation: Der Daten, der Forschenden, der Theorien und der Methodologien (Denzin 1978: 295ff). Der letztgenannte Typ ist jener, der in der vorliegenden Arbeit Anwendung findet, weshalb dieser Typus kurz genauer beschrieben werden soll. Methodologische Triangulationen können in zwei Formen unterschieden werden: »Within-methods« oder »between-methods«. Within-methods liegen dann vor, wenn multiple Strategien innerhalb einer Methode benutzt werden, um den Gegenstand zu untersuchen (ebd.: 301). »Between-methods«-Triangulationen sind wiederum jene, bei denen verschiedene Methoden zur Untersuchung des gleichen Gegenstands benutzt werden. Wie Denzin feststellt, sind hier oft die Schwachstellen der einen Methode die Stärken der anderen. In der Kombination beider, können die Forschenden so jeweils das Beste aus der jeweiligen Methode herausziehen, indem die Defizite der anderen gegenseitig überwunden werden (ebd.: 302). »Between-methods« können vielerlei Formen annehmen, aber ihre grundlegende Form besteht in der Kombination von zwei oder mehreren unterschiedlichen Forschungsstrategien, um die gleichen empirischen Einheiten zu untersuchen (ebd.: 302).

griffe im Diskurskorpus, kann in ihren qualitativen Verbindungen zu anderen Elementen, im zweiten Schritt so genauer nachgegangen werden.

Abbildung 4: Darstellung der Integration quantitativer und qualitativer Analyse zur Umsetzung der Diskurstheorie von Laclau und Mouffe auf der Designebene.



Quelle: Grafik angelehnt an Abb. 2 in Mayring 2001.

5.2 Lexikometrische Analyse des Diskurskorpus nach dem corpus-driven Ansatz

Im Folgenden wird zum einen die Lexikometrie als Methode vorgestellt und zum anderen der Ansatz veranschaulicht, mit dem mithilfe der Methode, das Diskurskorpus untersucht wird. Dabei wird der Ansatz der Korpusforschung von anderen abgegrenzt.

5.2.1 Lexikometrische Verfahren als Teil der Diskursanalyse

Allgemein betrachtet sind lexikometrische Verfahren² ein Instrument zur Untersuchung der quantitativen Beziehungen zwischen lexikalischen Elementen innerhalb geschlossener Textkorpora³ (Glasze 2013b: 101). Die Textkorpora müssen dabei in digitaler Form

2 Glasze (2013b: 101) verweist zudem darauf, dass Lexikometrie nicht der einzige verfügbare Begriff für Verfahren ist, die in diese Kategorie fallen. Seine Verwendung unterscheidet sich vor allem entlang der Sprache in denen über solcherart Verfahren gesprochen wird. So werden in der französischen Wissenschaftslandschaft hauptsächlich die Begriffe *léxicométrie* und *statistique textuelle* benutzt (Lebart/Salem/Berry 1998). In englisch- und deutschsprachigen Publikationen wird jedoch vielmehr von *corpus linguistics* bzw. *Korpuslinguistik* gesprochen (Teubert 1999, 2005).

3 Das Wort Korpus hat seinen Ursprung im lateinischen Wort *corpus* und steht dort für »Körper«. Gegenwärtig wird der Begriff hauptsächlich von den Sprachwissenschaften verwendet und bezeichnet im Allgemeinen laut Noah Bubenhofer eine »endliche Menge von konkreten sprachlichen Äußerungen, die als empirische Grundlage für sprachwissenschaftliche Untersuchungen dienen.« (Bubenhofer 2006b: o.S.)

vorliegen, was bedeutet, dass entweder auf digitale Texte zugegriffen wird oder die Texte mithilfe einer Variante der Texterkennung noch eingelesen werden müssen (Glasze 2008: 200). Als geschlossen können Textkorpora bezeichnet werden, wenn »deren Definition, Zusammenstellung und Abgrenzung klar definiert ist und die[se] nicht im Laufe der Untersuchung verändert werden.« (Glasze 2013b: 101) Wie Glasze betont, ist die Geschlossenheit insofern wichtig, als im Laufe des lexikometrischen Verfahrens unterschiedliche Teile der Textkorpora verglichen werden und nur dann sinnvolle Ergebnisse herauskommen können, wenn sie sich auf stabile oder anders gesagt auf »geschlossene« Zusammenstellungen von Texten beziehen (Glasze 2008: 200). Er macht deutlich, dass die Entscheidung für geschlossene Textkorpora zum einen auf eine inhärente Gefahr von Textanalysen zurück geht, die darin besteht, dass nur auf die Texte bzw. Textpassagen Bezug genommen wird, die den impliziten Erwartungen der Wissenschaftler*innen entsprechen. Die Arbeit mit geschlossenen Textkorpora beugt somit Risiken vor, in Zirkelschlüsse zu geraten. Zudem erhöhen sich damit die Chancen am Ende »Diskursmuster herausarbeiten zu können, die den impliziten Erwartungen [der Forschenden] zuwider laufen.« (Glasze 2008: 201) Der Fokus von Lexikometrischen Verfahren, die mit geschlossenen Korpora arbeiten, beruht damit auf einem transparenten »factually given discourse« und weniger auf Vorabinterpretationen der Forschenden (Glasze 2007: 664). Dies bedeutet jedoch nicht, dass damit überhaupt keine interpretativen Vorannahmen mehr herrschen würden. Bereits die Formulierung der Fragestellung sowie die Zusammenstellung der Textkorpora stellen bereits eine interpretatorische Leistung dar. Vielmehr wird lediglich der Schwerpunkt der Interpretation so weit wie möglich nach hinten in der Untersuchung verlagert, damit zuvor ausreichend Raum besteht, Anregung für mögliche unerwartete Erkenntnisse zu gewinnen (Glasze 2008: 199; Dzudzek 2013: 59). Dies wird umgesetzt, indem die Interpretation der Ergebnisse erst begonnen wird, nachdem die Ergebnisse der lexikometrischen Analyse bereits vorliegen (Glasze 2013b: 102). Um in der vorliegenden Arbeit dem Material die größtmögliche Chance einzuräumen, die Erwartungen des Forschenden zu verunsichern, um so auf möglicherweise unerwartete Zusammenhänge stoßen zu können, hat sich der Verfasser entschieden geschlossene Korpora zu verwenden.

Mit lexikometrischen Verfahren⁴ im Rahmen diskursorientierter Ansätze ergibt sich so die Möglichkeit, »Rückschlüsse auf diskursive Strukturen und deren Unterschiede zwischen verschiedenen Kontexten wie bspw. die Entwicklungen über die Zeit« oder zwischen institutionalisierten Gruppen ziehen zu können. Wie Glasze weiter schreibt ist es also das »Ziel lexikometrischer Verfahren in der Diskursforschung [...], großflächige Strukturen der Sinn- und Bedeutungskonstitution in Textkorpora zu erfassen.« (Glasze 2013b: 101) Lexikometrische Verfahren gehen dabei grundlegend davon aus, dass ihre Bedeutung und Produktion durch die »Materialität des Zeichens« erkennbar wird. Auf der materiellen Ebene eines Korpus können demnach Bedeutung, Kämpfe um Bedeutung und der Bedeutungswandel herausgearbeitet werden. Texte die aufgezeichnet sind,

4 Zunächst wurden lexikometrische Verfahren in den Sprachwissenschaften entwickelt. Konzeptionell beziehen sie sich dabei auf die Ausführungen der Linguistik von Ferdinand de Saussure und in Teilen auf jene Weiterentwicklungen im Poststrukturalismus wie auch die diskurstheoretischen Ideen Michel Foucaults (Glasze 2013b: 101f; Dzudzek 2013: 57).

sind damit als Träger von Bedeutung zu verstehen. Die kleinste Untersuchungseinheit mit der in lexikometrischen Verfahren Bedeutung nachvollzogen wird, stellen sogenannte »Textelemente«, also Wörter und nicht einzelne Wortzeichen wie Buchstaben, dar. Sie und ihr Verhältnis zum Gesamtkontext des Diskurses stellen die wesentliche zu untersuchende Materialität in lexikometrischen Verfahren dar (Dzudzek 2013: 57). So besteht mit ihnen ein Werkzeug, mit dem »die Unterschiedlichkeit von Verweisstrukturen und damit der Bedeutungen von einzelnen Wörtern und Zeichenverkettungen in unterschiedlichen diskursiven Formationen« erfasst werden kann (Glasze 2013b: 102).

Wie Derya Gür-Şeker (2014: 599) darlegt, gibt es innerhalb lexikometrischer Verfahren zwei Wege, wie die Analysen von Textkorpora stattfinden können. Beide Verfahren werden auch innerhalb von korpusorientierten Diskursanalysen angewendet: Zum einen gibt es den deduktiven Ansatz, der auch als »*corpus-based*« oder »korpusbasiert« bezeichnet wird. Hiermit werden korpuslinguistische Ansätze beschrieben, die an die Textkorpora mit einer bereits vordefinierten Hypothese herantreten. Damit die Forschenden ihr Interesse, »eine spezifische Hypothese bezüglich des Untersuchungsgegenstandes mittels Korpusdaten« überprüfen können, wird das Korpus nach festgelegten Kategorien untersucht (ebd.: 599). Wie Noah Bubenhofer aber kritisch anmerkt, besteht dabei immer die Gefahr, am Ende in den Daten nur die Strukturen zu finden, die mit der Theorie im Einklang stehen, mit der die Forschenden an das Material herangegangen sind. All das, was quer zu einer Theorie steht, ist damit nicht zu entdecken (Bubenhofer 2009: 101).

Demgegenüber steht die induktive Herangehensweise, welche auch als »*corpus-driven*« oder »korpus-gesteuert« bezeichnet wird. Hier werden die »im Korpus vorliegenden Daten zunächst beobachtet, um anschließend aus dieser Beobachtung heraus Regeln für spezifische sprachliche Phänomene abzuleiten.« (Gür-Şeker 2014: 599) Durch das induktive Vorgehen, so Bubenhofer, wird das Korpus nicht voreilig auf der Basis der theoretischen Vorannahmen analysiert, sondern die ersten Schritte der lexikalischen Kategoriebildung werden strikt aus dem Korpus selbst generiert. Erst nach der Analyse erfolgt die erste Hypothesenbildung über die Daten (Bubenhofer 2009: 17). Das Korpus dient demnach als »Inspirationsquelle« bezüglich der Entwicklung von Hypothesen über die darin auftauchenden lexikalischen Phänomene und Strukturen (Gür-Şeker 2014: 599). Methodisch gesehen heißt das, ohne im Voraus vorgefertigte Suchanfragen an das Korpus heranzugehen, um dabei die Chance zu haben auch »auf Strukturen zu stoßen, an die man nicht schon vor der Untersuchung gedacht hat.« (Glasze 2013b: 102) Vor allem der explorative Charakter des »*corpus-driven*« Ansatzes eignet sich besonders, »um einen ersten Überblick über Unterschiede und Gemeinsamkeiten sprachlicher Verweisstrukturen« zu gewinnen (ebd.: 102). Da die vorliegende Arbeit den Anspruch hat die Interpretation so weit wie möglich nach hinten zu verlagern, um genügend Raum für explorative Überraschungen im Diskurskorpus zu finden, wurde sich für den »*corpus-driven*« Ansatz entschieden.

Gemäß Glasze lassen sich lexikometrische Verfahren im Anschluss an die Diskurstheorie nach Laclau und Mouffe als ein Teilschritt der Operationalisierung verwenden, da sie dazu beitragen können, den Status von lexikalischen Elementen und ihre Beziehungen zu anderen lexikalischen Elementen zu untersuchen, um dadurch Strukturen der Bedeutungskonstitution herauszuarbeiten. Die Untersuchung von Kontexten lexi-

kalischer Elemente steht somit einer Grundannahme der Diskurstheorie sehr nahe, da hier »Bedeutung als ein Effekt der Beziehung von (lexikalischen) Elementen zu anderen (lexikalischen) Elementen« begriffen wird (Glasze 2013b: 102). Der Einsatz von lexikometrischen Verfahren in Diskursanalysen geschieht bisher eher selten, vor allem in der deutschsprachigen und englischsprachigen diskursanalytischen Landschaft. Vielfach bleibt ihre Verwendung auf die Bereiche der Linguistik beschränkt. Sie erfahren deshalb nur selten eine Anwendung in gesellschaftspolitischen Forschungsprojekten (Baker 2006: 1ff; Glasze 2013b: 102f). Die vorliegende Arbeit stellt einen Beitrag dazu dar, lexikometrische Verfahren im Rahmen einer politischen Diskursanalyse in der Disziplin der IVE zur Anwendung zu bringen, um so das transnational gesellschaftspolitisch äußerst relevante Feld Globaler Bildung zu untersuchen.

5.2.2 GB in Global Governance-Strukturen als thematisch orientierter Diskurskorpus

In den methodologischen Diskussionen zur Korpuslinguistik wird angemerkt, dass bereits ein Korpus⁵ mit zehn- und zwanzigtausend Textwörtern eine verlässliche Auskunft über die zu untersuchende Fragestellung geben kann (Scherer 2014: 7). Folglich wurde in der Zusammenstellung der Textkorpora für die lexikometrischen Untersuchungen darauf geachtet, dass jedes Textkorpus der Fallstudie zunächst jeweils eine spezifische Sprecher*innenposition der Global Governance-Architektur im Bildungsbereich repräsentiert. Darüber hinaus wurde sicher gestellt, dass jedes Textkorpus aus mindestens zehntausend Wörtern besteht, um ihn als aussagekräftigen Untersuchungsgegenstand werten zu können (siehe Tab. 2 unten).

Im Rückbezug auf den Diskursforschenden Siegfried Jäger muss bereits die Art und Methode, wie das Korpus zusammengestellt ist, als erster wichtiger Forschungsschritt angesehen werden (Jäger 1997). Allein der Akt die geeigneten Sprecher*innenpositionen des zu untersuchenden Feldes in Wechselwirkung mit der eigenen Forschungsfrage ausfindig zu machen sowie die Suche nach einem homogenen Textgenre der jeweiligen Sprecher*innenpositionen stellt bereits einen großen Schritt der eigentlichen Forschung dar. Auch die daran anschließende Arbeit der Aufbereitung digitaler Texte für die computergestützte Analyse darf in ihrem Arbeitsaufwand nicht unterschätzt werden.

Die Frage nach dem Aufbau eines eigenen Korpus wurde zu Beginn der eigentlichen Analyse gestellt. Der Aufbau eines eigenen Diskurskorpus war notwendig, da es bisher zum Thema keinen öffentlich verfügbaren Referenzkorpus gibt, um die Frage nach der diskursiven Formation GB in Global Governance-Strukturen zu beantworten. Bevor also in diesem Falle mit einer Diskursanalyse begonnen werden konnte, musste ein eigenes Korpus als Untersuchungsgegenstand, das der Forschungsfrage gerecht wird, aufgebaut

5 Wenn wir nachfolgend von einem Korpus sprechen, wird damit eine Sammlung von Texten oder Textteilen, die nach spezifischen wissenschaftlichen Kriterien ausgewählt und geordnet wurden, bezeichnet (Scherer 2014: 3). In der hier vorliegenden Untersuchung handelt es sich um Texte, die virtuell zu Verfügung standen. Zumeist als PDF oder in Einzelfällen als Websitebasiertes Produkt der Schriftsprache.

werden (Gür-Şeker 2014: 586). Der Aufbau eines eigenen Korpus als Untersuchungsgegenstand orientierte sich an der Forschungsfrage und dem Erkenntnisinteresse, weshalb modellhafte Sprecher*innenpositionen – verschiedenste institutionalisierte Gruppen und Diskurskoalitionen mit transnationalem Wirken – der Global Governance-Architektur im Bildungsbereich zu einzelnen Textkorpora aufbereitet wurden, um letztendlich ein geschlossenes Gesamtdiskurskorporus zu bilden. Wie Glasze verdeutlicht, geht es vor der Zusammenstellung des Diskurskorporus darum, Überlegungen anzustellen, bezüglich welcher Kriterien die Bedeutungskonstitutionen untersucht und verglichen werden sollen. Wie er weiter ausführt, hat dies Auswirkungen darauf, wie das Korpus aufgeteilt bzw. segmentiert wird (Glasze 2013b: 104).

In der vorliegenden Diskursanalyse geht es darum, die Bedeutungskonstitutionen unterschiedlicher Sprecher*innenpositionen der GB und ihre Macht über Inklusion und Exklusion im Zuge ihrer globalen Bildungsbestrebungen in einem festgelegten Abschnitt zu Beginn des 21. Jahrhunderts zu untersuchen. Laut Glasze ist es wichtig bei einem sich wechselnden Merkmal auf dem der Fokus liegt, die anderen Merkmale der Texte konstant zu halten (ebd.: 104f). Deshalb wurde bei der Auswahl der Texte darauf geachtet, dass es sich bei der Zusammenstellung um ein zeitlich homogenes Korpus handelt, zusammengesetzt aus Texten unterschiedlicher Sprecher*innenpositionen im Bereich GB aus unterschiedlichen Ebenen der Architektur der Global Governance-Architektur. Jedes Teilkorpus des Diskurskorporus setzt sich wiederum aus einer homogenen Sprecher*innenposition zusammen.

Die Zusammenstellung der ausgewählten Sprecher*innenposition stellt eine höchstmögliche Vielfalt der Handlungsebenen der inter- und transnational agierenden Akteur*innen innerhalb der Global Governance-Architektur dar. Die Vielfalt der hier vorgestellten Handlungsebenen GB orientiert sich an den allgemeinen Ausführungen Dirk Messners zur Architektur der Global Governance (Messner 1999: 13), die als ein Modell gelesen werden kann, welches als Vorlage zur Auswahl bildungsgestalterischer Akteur*innen im globalen Maßstab für das Diskurskorporus angesehen werden kann. Im Anschluss an dessen Ausführungen wurden Bildungsakteur*innen, die auf den jeweiligen Ebenen agieren – von der Ebene der UN, der Global Private Player bis hin zur internationalen und lokalen Zivilgesellschaft – ausgewählt. Diese stehen exemplarisch für mindestens eine hier diskursrelevante Sprecher*innenposition der jeweiligen Ebene in der Architektur der Global Governance. Mit der Integration mindestens einer diskursrelevanten Sprecher*innenposition von vier verschiedenen Handlungsebenen soll die Repräsentativität verschiedener Handlungsebenen modellhaft in dem untersuchten Diskurskorporus gewährleistet werden. Dies geht auf die Anmerkungen Gür-Şekers zurück, die betont, dass ein Diskurskorporus repräsentativ ist, wenn »weder wesentliche Diskurskomponenten fehlen, noch [...] bestimmte Komponenten überbetont werden.« (Gür-Şeker 2014: 593) Kriterien dafür sind, dass diskursrelevante Texte das Korpus prägen sowie eine Vielfalt an Sprecher*innenpositionen vorhanden ist (ebd.: 593).

Es ist somit versucht worden einen modellhaften synchronen⁶ diskursanalytischen Schnitt durch die Architektur der neuen globalen bildungsgestalterischen Arena zu ziehen. Dass diese Auswahl in Rückbezug auf Siegfried Jägers Anmerkungen zur Durchführung von Diskursanalysen nicht erschöpfend sein kann, ergibt sich aus seiner Anmerkung zur Gewinnung verlässlichen Materials (Jäger 2012: 123). Er führt aus, dass »das gesamtgesellschaftliche Archiv [...] ›in seiner Totalität nicht beschreibbar‹ [...]« ist, sondern es geht vielmehr darum »dieses Geflecht analytisch zu entwirren und in all seinen Verästelungen genau zu beschreiben«, um es danach untersuchen zu können (ebd.: 123). Dies, wie er hervorhebt, »ist weder für eine bestimmte aktuelle Gesellschaft insgesamt möglich und erst recht nicht für diachrone Gesellschaftsanalysen« (ebd.: 123). Obwohl es sich in dem vorliegenden Forschungsvorhaben um eine synchrone Analyse zu Beginn des 21. Jahrhunderts handelt, erhebt die hier getroffene Auswahl nicht den Anspruch auf Vollständigkeit, geschweige denn den Anspruch auf die Totalität seiner Aussagen. Die Begründungen für die genauere Auswahl der spezifischen Sprecher*innenpositionen als analytischer Zugang zum Diskurs werden im Kapitel 6.3 zum Forschungsdesign genauer vorgestellt.

Mit der Untersuchung sollen Erkenntnisse erzielt werden können, die darüber Aufschluss geben, wie im Diskurs der GB mit den wesentlichen Herausforderungen der gegenwärtigen Globalisierung – der wachsenden globalen Ungleichheit und der zunehmenden ökologischen Verwüstung – umgegangen wird. Damit, wie Jäger anmerkt, Diskursanalysen in Bezug auf die Materialfülle bewältigbar bleiben, wird sich, wie in der hier vorliegenden Arbeit der Fall, »auf einen stark eingegrenzten Problembereich und einen überschaubaren historischen Zeitraum« konzentriert (Jäger 2012: 126). Der zeitliche Ausgangspunkt aller ausgewählten Dokumente der lexikometrischen Analyse setzt deshalb im Jahr 2012 an und erstreckt sich zur Erfassung des Diskurses bis zum Jahr 2018. Der angesetzte Zeitraum von 2012 bis 2018 ist als ein Ausschnitt zu begreifen, in dem die historisch-politischen Herausforderungen, die durch die in Kapitel 2 beschriebene postkolonial-neoliberale Globalisierung zu Beginn des 21. Jahrhunderts entstanden sind, bereits in die diskursive Formation internationaler Gestaltung GB eingeflossen sind und sich damit die Auseinandersetzung mit den gegenwärtigen Herausforderungen einer wachsenden globalen Ungleichheit und der globalen ökologischen Verwüstung im Diskurs GB bereits untersuchen lässt.

Bei der Zusammenstellung des hier untersuchten Diskurskorpus im Arbeitsfeld Globaler Bildung muss nach Gür-Şeker in Anlehnung an Teubert von einem *thematisch orientierten Korpus* gesprochen werden, für das einige wichtige Grundlagen erfüllt sein müssen (Gür-Şeker 2014: 585f). Dazu wurde durch die Festlegung bestimmter Parameter ein Diskurs der GB anhand eines Korpus konstituiert. Die für die hier vorliegende Forschung bedeutsamen Parameter waren das Thema, der Zeitraum, das Areal, die

6 Unter einer synchronen Diskursanalyse wird nach Margarete Jäger und Siegfried Jäger diejenige verstanden, die einen »synchronen Schnitt« durch den Diskursstrang zieht und damit eine »gewisse (endliche) Bandbreite« hat. Solch ein Schnitt bezieht sich auf die Sagbarkeiten zu einem bestimmten »Zeitpunkt«. Wie sie betonen, kann auch eine »Zeitdauer« als »Zeitpunkt« fungieren, wenn dieser z.B. die Dauer eines historisch wichtigen Ereignisses oder einer politischen Auseinandersetzung umfasst (Jäger/Jäger 2007: 26).

Veröffentlichungsform sowie die intertextuelle Beziehung dieser Texte, des hier vorgestellten Diskurskorpus, zueinander. Nachfolgend wird spezifiziert, wie diese einzelnen Parameter in dem hier vorliegenden Korpus angewandt wurden:

- Thema: Analyse der diskursiven Formation Globaler Bildung in den Global Governance-Strukturen insbesondere hinsichtlich der Verhandlung der Themen globale Ungleichheit und ökologische Verwüstung
- Zeitraum: 2012–2018
- Areal: Querschnitt durch die Architektur der Global Governance im Bereich Globaler Bildung auf der didaktischen Ebene internationaler Bildungsgestaltung und -legitimation
- Veröffentlichungsform: Strategie- und Konzeptpapiere, Berichte, Programme sowie bildungstheoretische und -praktische Beiträge des Bereichs als PDF-Dokumente
- Intertextuelle Beziehung: Strategie- und Konzeptpapiere im Feld Globaler Bildung, die zusammengekommen, ob bewusst oder unbewusst, an der Gestaltung und Legitimation spezifischer Artikulationen Globaler Bildung arbeiten

An den Ausführungen Gür-Şekers orientiert beansprucht die Erstellung dieses thematisch orientierten Korpus nicht, dass damit die Vollständigkeit des Diskurses der GB abgebildet wird. Vielmehr geschieht dies mit dem Bewusstsein, dass es sich dabei nur um einen Diskursauschnitt handelt und der Diskurs sich darin nicht erschöpft (Gür-Şeker 2014: 586). Im Rückgriff auf Fritz Herrmanns spricht Gür-Şeker auch von einem »konkreten Korpus«, was bedeutet, dass solcherlei Diskursanalysen lediglich auf jener Auswahl von Daten basieren, die der Analyse zugrundeliegen. Damit der Aufbau begonnen werden konnte, mussten zudem noch einige Fragen wie die »Verfügbarkeit der Texte bzw. Daten, Textsorte, Größe der Texte und Umfang der gespeicherten Elemente, d.h. Textfragmente oder komplette Texte« geklärt werden (ebd.: 590):

- Verfügbarkeit: Auf der Ebene der Verfügbarkeit, handelt es sich um digitale Dokumente, die im Internet auf der Website der ausgewählten Organisationen, die hier als Sprecherpositionen den Ausgangspunkt für die lexikometrischen Untersuchungen bildeten, verfügbar waren. Es konnte damit auf bereits elektronisch verfügbare Daten zurückgegriffen werden.⁷
- Format: Das Format der für die diskursanalytische Untersuchung ausgewählten Quellen ist hauptsächlich ein PDF-Format. Eine Ausnahme stellen die beiden Ausgaben des Jahres 2012 der Zeitschrift Adult Education and Development dar, welche nur als Webtext vorhanden waren, sodass die Texte in eine TXT-Datei überführt wurden (DVV-International 2012a, 2012b).
- Textsorte: Bei den ausgewählten Materialien für die lexikometrische Untersuchung handelt es sich im Hinblick auf die Textsorte zumeist um Strategie- und Konzeptpa-

7 Wie Noah Bubenhofer bereits festgestellt hat, hat das World Wide Web zum Bau eines eigenen Korpus in lexikometrisch oder korpuslinguistischen Verfahren immer mehr an Bedeutung gewonnen, da hier der Möglichkeit an Quellen für das Korpus zu kommen, keine Grenzen gesetzt sind (Bubenhofer 2006a: o.S.).

piere, Berichte und Programme sowie bildungstheoretische und -praktische Beiträge des Bereichs GB der jeweiligen Sprecher*innenpositionen, die die von ihnen vertretene organisationsphilosophische Ideologie ihrer Bildungsbestrebungen darstellen. Ebenso hielten auch zwei Zeitschriften Einzug in die lexikometrischen Untersuchung, da sie die Textsorte sind, mit der auf der Ebene der Bezugs-Wissenschaft in den Diskurs der GB interveniert wird. Sie alle stellen Dokumente dar, mit denen analysiert werden kann, wie die unterschiedlichen Sprecher*innenpositionen ihre Vorstellungen von GB als Antwort auf die Herausforderungen der gegenwärtigen Globalisierung artikulieren.

- **Umfang der Texte:** Zur lexikometrischen Untersuchung der Texte wurde sich auf die Hauptteile der Dokumente konzentriert, d.h. es wurde alles entfernt, was nicht zum eigentlichen Text gehört. Iris Dzudzek, die sich ebenso lexikometrischer Verfahren bedient, zählt dazu das Inhaltsverzeichnis, alle Kopfzeilen, Zahlen, alphabetische Nummerierungen und einzeln stehende Buchstaben (Dzudzek 2013: 74). Weiterhin wurden Grafiken, Tabellen und Koordinatensysteme entfernt, da diese nicht in Textstrukturen überführt werden konnten, aus denen sich in der Umwandlung ins TXT-Format, die zuvor im gelayouteten Dokument vorhandene Sinnstruktur für lexikometrische Untersuchungen hätte erschließen können. Lexikometrische Programme können in der Regel nur Wörter auswerten und keine Grafiken.

Genaue Informationen über die Zusammensetzung der Korpora zur lexikometrischen Untersuchung im Rahmen der hier vorgelegten Fallstudie, welche Größe sie jeweils umfassen und welche spezifischen Sprecherpositionen zu einem Teilkorpus zusammengefasst worden sind, wird in Kapitel 6.3 dargelegt.

5.2.3 Methoden lexikometrischen Arbeitens mithilfe des Programms Lexico 5

Lexico⁸ ist ein Programm für lexikometrische Verfahren, welches zur automatisierten Auswertung von Textkorpora verwendet wird. Es stellt verschiedenste Funktionen zur Lexikometrischen Analyse bereit. So lassen sich Frequenz-, Konkordanz- und Kookkurrenzanalysen sowie Analysen zu den Charakteristika eines Teilkorpus durchführen (Fracchiolla et al. 2004: 5; Dzudzek 2013: 74). Darüber hinaus kann das Programm auch grundlegende Aussagen zu einem Korpus ermitteln, wie z.B. die absolute Anzahl

8 Lexico wurde von Cédric Lamalle, William Martinez, Serge Fleury und André Salem an der »Université de Sorbonne nouvelle – Paris 3« entwickelt. Die erste Version von Lexico wurde 1990 veröffentlicht. Seitdem wurde das Programm konstant weiterentwickelt (Fracchiolla et al. 2004: 5). Gegenwärtig ist bereits die fünfte Version des Programms verfügbar, welche auf der Website Lexicos heruntergeladen werden kann. Für das Programm gibt es unterschiedliche Nutzungsbestimmungen. In den meisten Fällen muss für die Nutzung des Programms eine Lizenz erworben werden, jedoch steht das Programm für Studierende oder Wissenschaftler*innen ohne institutionelle Anbindung zur kostenfreien Nutzung zur Verfügung. Der Wortlaut der Nutzungsbestimmungen dazu im Original: *Règles d'utilisation de Lexico 3.6 et de Lexico 5.beta: Si vous êtes un étudiant ou un chercheur isolé : vous pouvez télécharger les logiciels et les utiliser gratuitement pour votre travail. Notez que si un jour une des institutions à laquelle vous appartenez décidait d'acquérir une ou des licences, cela nous aiderait considérablement à développer nos logiciels.*

der Okkurrenzen,⁹ die helfen die Größe eines Korpus zu bestimmen, die Anzahl der unterschiedlich benutzten Wörter sowie *Hapax Legomena*¹⁰ (Dzudzek 2013: 74).

Für die computergestützte korpuslinguistische Untersuchung mussten die Texte aufbereitet werden. Dazu wurden die unterschiedlichen Dateiformate der Texte, zumeist PDF- und DOC-Dateien, computergestützt¹¹ in TXT-Formate überführt (Fracchiolla et al. 2004: 9). Für die lexikometrische Analyse wurde sich für die Benutzung des Programmes *Lexico 5*¹² entschieden. Über die Formatierung in das TXT-Format hinaus verlangt das Programm *Lexico 5* einen weiteren Formatierungsschritt der Texte. In lexikometrischen Studien werden Frequenzen verschiedener Textformen in verschiedenen Teilen des Korpus miteinander verglichen. Damit solch ein Vergleich möglich wird, müssen die einzelnen Texte mit *tags*¹³ versehen werden, sodass die logische Struktur des Korpus durch das Programm nachvollzogen werden kann (Fracchiolla et al. 2004: 11). Darüber hinaus wurden die einzelnen Absätze im Text markiert. Dies dient dazu die relativen Distanzen zwischen den Wörtern und Wortfolgen in großen Textmengen genauer auslesen zu können. Die Markierungen wurden für *Lexico* mithilfe des Zeichens »\$« am Beginn jedes Absatzes gesetzt. Dieses stellt die vorgegebene Key-Form für *Lexico* dar. (Fracchiolla et al. 2004: 12) Die gesamten Textdokumente wurden zudem in Kleinschreibung umgewandelt. Wie Dzudek betont, kann so sichergestellt werden, dass gleiche Wörter, einmal in Groß- und in Kleinschreibung nicht differenziert betrachtet werden. Die großgeschriebenen Wörter von den kleingeschriebenen Wörtern im Korpus zu unterscheiden wird deshalb nicht als sinnvoll erachtet (Dzudzek 2013: 75).

5.2.3.1 Frequenzanalyse

Um sich aus synchroner Perspektive dem Diskurs der GB in den Global Governance-Strukturen anzunähern, wurden Frequenzanalysen mit den Teilkorpora der jeweiligen Sprecher*innenpositionen der Architektur der Global Governance durchgeführt, so dass mögliche Teilformationen des Diskurses entlang unterschiedlicher Sprecher*innenpositionen aufgedeckt werden können. Frequenzanalysen bieten die Möglichkeit absolute

-
- 9 »Okkurrenzen bezeichnet das Vorkommen einer sprachlichen Form wie Wörter oder Satzzeichen.« (Glasze 2013b: 130)
 - 10 Dies sind Textelemente, die nur einmal im Text vorkommen, womit sie einen Hinweis über Wortneubildungen im Korpus ausdrücken.
 - 11 Hierzu bieten sich Websites an, die diese Funktion bis zu einer bestimmten Größe kostenfrei übernehmen. Es ist jedoch mit Vorsicht vorzugehen, da diese zum größten Teil nicht fehlerfrei funktionieren.
 - 12 Auch wenn noch weitere mögliche Programme zur korpuslinguistischen Analyse bereit stehen, so ist der Vorteil dieses Programmes, dass es für Individualforschende kostenfrei zur Verfügung steht und zum anderen noch wichtigeren Punkt als einziges korpuslinguistisches Programm die Möglichkeit hat, auch *segment repetes* (auch *N-gramme* genannt) zu analysieren. Dies ist von Vorteil, da in der vorliegenden Analyse des Diskurses Wortverbindungen wie »sustainable developement« oder »global citizenship« eine wichtige Rolle spielten. Die Möglichkeit nur die Häufigkeit einzelner Worte nachvollziehen zu können, würde dem Anspruch der Forschung nicht gerecht werden.
 - 13 Für die Aufbereitung der Texte wurde sich für die *tags* Sprecher*innenposition, Jahr und Nummer der Ausgabe entschieden. Die dadurch entstandene Textgrammatik sieht beispielhaft wie folgt aus: Sprecherposition UNESCO = <actor=UNESCO>/Jahr 2015= <year=2015>/81. Nummer der Ausgabe = <number=81>

oder relative Häufigkeiten charakteristischer Wortformen in einem bestimmten Teilkorpus nachvollziehen zu können. Unter charakteristischen Wortformen werden bei Lexico einzelne Wörter, aber auch sogenannte *repeated segments* verstanden. *Repeated segments* sind häufig miteinander auftauchende Wortfolgen (Fracchiolla et al. 2004: 21). Es können so Hinweise für die Feinanalyse gefunden werden. Lexico ermöglicht es die Frequenzen von Textelementen in absoluten oder relativen Häufigkeiten eines Korpus zu errechnen. In dem hier vorliegenden Verfahren wurde sich für die Betrachtung der relativen Häufigkeit entschieden, da hier die Größe des Teilkorpus nicht zum entscheidenden Kriterium der Frequenz gemacht werden soll. Bei der Analyse absoluter Häufigkeiten wäre das Problem, dass ein Teilkorpus mit großer Textmenge historisch bedeutender erscheinen würde als kürzere Texte (Dzudek 2013: 75). Frequenzanalysen können aber nur eine erste Annäherung an mögliche Teilformationen des Korpus darstellen. Die Häufigkeit eines Textelements in einem Korpus oder Teilkorpus ist noch kein valider Beweis für seine Bedeutung für den Diskurs, wie Dzudek betont (ebd.: 75). Um die Relevanz der Textelemente für den Diskurs näher zu bestimmen wird eine weitere formale Methode herangezogen.

5.2.3.2 Characteristic Element Diagnostic

Eine weitere wichtige Methode zur Bestimmung relevanter Textelemente stellt die Untersuchung der »Spezifität der Verknüpfung lexikalischer Elemente« dar (Glasze 2008: 200), nach Lebart, Salem und Berry auch als Characteristic Element Diagnostic, kurz *ced* bezeichnet (Lebart/Salem/Berry 1998: 129). Es handelt sich dabei um eine Methode, die nach Lebart/Salem/Berry aus der klassischen statistischen Analyse für die Textanalyse adaptiert wurde. Mithilfe von Characteristic Element Diagnostic (CED) kann gezeigt werden, welche lexikalischen Formen in einem Teil des Korpus verglichen mit dem Gesamtkorpus bzw. mit einem anderen Teilkorpus besonders charakteristisch sind (Lebart/Salem/Berry 1998: 129; Glasze 2013b: 107). Die Charakteristik, die zugleich ein Ausdruck für die Bedeutsamkeit eines Textelements in einem bestimmten Teilkorpus sein kann, geht dabei auf das Maß der Wahrscheinlichkeit zurück, mit der ein lexikalisches Element in einem Teilkorpus auftritt. Die Methode der CED eignet sich deshalb z. B. dazu festzustellen, wie spezifisch ein lexikalisches Element für einen Teilkorpus ist und diese mit anderen Teilkorpora zu vergleichen. Dadurch kann festgestellt werden, ob die Begriffe von allen Teilkorpora geteilt werden oder spezifisch für einen bestimmten Teil des Diskurskorpus sind. Damit können Hinweise zu möglichen Knotenpunkten oder über von bestimmten Sprecher*innenpositionen häufig verwendete Elemente gewonnen werden.

Gemäß Dzudek (Dzudek 2013: 75f), die sich auf Lebart et al. (Lebart/Salem/Berry 1998: 130–136) bezieht, liegt dem Verfahren der CED eine Wahrscheinlichkeitsfunktion zugrunde. Sie dient dazu, zu bestimmen, ob die Textelemente, die in einem Textabschnitt vorkommen, charakteristisch für diesen Textabschnitt sind. Zur Ermittlung der Charakteristik wird »aus der absoluten Häufigkeit eines Textelements im Gesamtkorpus, der Gesamtzahl aller im Korpus vorkommenden Textelemente sowie der Gesamtzahl aller in einem Teilkorpus vorkommenden Textelemente die Wahrscheinlichkeit bestimmt, mit der ein bestimmtes Textelement in einem Teilkorpus vorkommt.« (Dzudek 2013: 75f)

Lexico 5 ermittelt die Charakteristik computerbasiert in dem »aus dem Verhältnis zwischen der Häufigkeit des Wortes und der Gesamtzahl aller Wörter im Korpus die Wahrscheinlichkeit für eine bestimmte Frequenz des Wortes in einem Teil des Korpus berechnet«¹⁴ wird (Glasze 2013b: 110). Die Charakteristik des Vorkommens eines bestimmten Textelements wird bei Lexico 5 als »negativer Exponent der Zehnerpotenzen dieser Wahrscheinlichkeiten« angegeben. Diese treten in der Form 10^{-x} auf und werden bei Lexico als Spezifität bezeichnet¹⁵ (ebd.: 110). Der Exponent ist ein Ausdruck für »die Größe der Abweichung des charakteristischen Wortes von einer zufälligen Verteilung des Wortes im Korpus« (Dzudek 2013: 76). Mithilfe der Spezifitäten lassen sich somit Aussagen darüber treffen, »welche Wörter bzw. Wortfolgen in einem Teilkorpus im Vergleich zum Gesamtkorpus spezifisch häufiger bzw. seltener vorkommen.« (Glasze 2013b: 110)

Bei CED wird nur der Exponent als Maß der Spezifität aufgeführt (Lebart/Salem/Berry 1998: 134; Dzudek 2013: 76). Nun gibt der Exponent grundlegend nur an, ob ein Begriff charakteristisch für einen Textabschnitt ist oder nicht. Es gibt jedoch zwei unterschiedliche Varianten inwiefern der Gebrauch eines Begriffs in einem Abschnitt charakteristisch sein kann. So besteht die Möglichkeit zum einen der Über- und zum anderen der Unterrepräsentation des Wortes oder der Wortfolgen in einem Abschnitt. Zur Unterscheidbarkeit beider Fälle wurde die CED um eine Möglichkeit der Unterscheidung erweitert. Die Unterscheidung wird durch ein Plus- bzw. Minuszeichen vor dem Exponenten ausgedrückt. Laut Dzudek (2013: 77), gibt ein Pluszeichen vor dem Exponenten ein »positiv charakteristisches Textelement« an. Hiermit wird zum Ausdruck gebracht,

-
- 14 Laut Lebart/Salem/Berry (1998: 130 – 136) wird für die Berechnung auf einen Wahrscheinlichkeitstest unter einer hypergeometrischen Verteilung mit den Variablen »Gesamtzahl der Okkurrenzen des Gesamtkorpus«, »Okkurrenzen des bestimmten Wortes bzw. der Wortfolge im Gesamtkorpus« und »Gesamtzahl der Okkurrenzen im Teilkorpus« zurückgegriffen (Glasze 2008: 202). Dazu wird errechnet, wie wahrscheinlich der Fall ist, dass wenn ein bestimmtes Textelement y-mal im Gesamtkorpus und x-mal im Teilkorpus vorkommt, »genau diese Situation bei einem Wahrscheinlichkeitstest unter einer hypergeometrischen Verteilung eintritt.« (Dzudek 2013: 76) Die Grundannahme so Dzudek weiter, die dahinter steht ist, dass wenn der dabei herauskommende »Wahrscheinlichkeitswert mittels einer hypergeometrischen Verteilung beschrieben werden [kann], [...] das Auftauchen des Textelements in dem Teilkorpus zufällig ist. Weicht die Wahrscheinlichkeit von der Funktion [der hypergeometrischen Verteilung jedoch] signifikant ab, so ist das Auftreten des Textelements für das entsprechende Teilkorpus charakteristisch.« (ebd.: 76) Folglich kann festgehalten werden, dass je größer die Wahrscheinlichkeit ist, desto eher das Auftauchen des Textelements als zufällig bewertet werden muss. Eine Feststellung die andersherum ebenso gilt: Je kleiner die Wahrscheinlichkeit, desto eher die Annahme, dass das Textelement nicht zufällig im Teilkorpus auftaucht und deshalb als äußerst charakteristisch für das Teilkorpus eingestuft werden muss (ebd.: 76).
- 15 Glasze (2008: 202) hebt hervor, dass in den Sozialwissenschaften vielfach mit einem Signifikanzniveau von 5 % oder 1 % gearbeitet wird, was bedeuten würde, dass ein charakteristisches Textelement 5 bzw. 1 mal in 100 Textelementen auftreten müsste. Allerdings ist solch ein signifikantes Vorkommen von Textelementen aufgrund der Größe des allgemein verwendeten Wortschatzes unrealistisch (s. auch Dzudek 2013: 76). Die empirisch gefundenen Wahrscheinlichkeiten sind gemäß Glasze bei lexikometrischen Verfahren, zumeist um ein Vielfaches kleiner als 10^{-5} . Dzudek (2013: 76) fügt dem ganzen deshalb hinzu, dass im Rahmen von Characteristic Element Diagnostics nur der Exponent als Maß der Spezifität angegeben wird.

dass es sich dabei um ein Textelement handelt, welches in dem entsprechenden Teilkorpus besonders häufig vorkommt und demnach als explizit charakteristisch zu deuten ist. Im Gegensatz dazu weist ein Minuszeichen vor dem Exponenten auf ein »negativ charakteristisches Textelement« hin, was wiederum darauf verweist, dass das Textelement in dem entsprechenden Teilkorpus ausgesprochen selten auftritt. Mit dem Zeichen + oder – vor dem Exponenten wird demnach, wie Dzudek noch einmal zusammenfasst, der »über- bzw. unterdurchschnittliche Gebrauch eines Wortes in einem Teilkorpus im Verhältnis zum Gesamtkorpus« angezeigt (Dzudzek 2013: 77).

In der hier vorliegenden Analyse werden ausschließlich »positiv charakteristische Textelemente« in die Untersuchung einbezogen. Dies wird damit begründet, dass eine relative Überrepräsentation eines Begriffes in einem bestimmten Abschnitt eher darauf schließen lässt, dass es sich um einen Begriff handelt, der die Funktion eines den Diskurs charakterisierenden Textelements in dem entsprechenden Abschnitt übernimmt. Im Sinne von Laclau und Mouffe kann damit auch eher davon ausgegangen werden, dass es sich dabei um einen charakteristischen Begriff handelt, der durch seinen häufigen Gebrauch in die hegemoniale Ordnung als Knotenpunkt mit eingeflossen ist (Glasze 2008: 203). Darüber hinaus ist für lexikometrische Verfahren in Diskursanalysen noch kein festgelegtes Signifikanzniveau bestimmt worden, d.h. ein Niveau, ab dem es sich verbindlich um ein charakteristisches Element handelt. Im Rahmen der vorliegenden Arbeit wurde das Signifikanzniveau auf mindestens 10^{-10} bzw. $X \geq 10$ festgelegt (orientiert an: Glasze 2008: 202). Damit sollte sichergestellt werden, dass zum einen genügend Elemente im Vorhinein aussortiert werden, um eine adäquate Feinanalyse zu ermöglichen und zum anderen trotzdem noch ausreichend Elemente zur Verfügung stehen, um eine anspruchsvolle Analyse zu ermöglichen. Ein weiteres Kriterium, welches in die Auswahl charakteristischer Elemente einfließt, auch wenn sie sich oberhalb des festgelegten Signifikanzniveaus befinden, ist, dass es sich bei den ausgewählten nicht um beliebige Elemente handelt. Vielmehr ist entscheidend, dass es Elemente sind, die keine reine Selbstreferentialität der Sprecher*innenposition, die untersucht wurde, darstellen, sondern inhaltlich auf das Phänomen der GB bezogen sind.

5.2.3.3 Kookkurrenzanalysen

Für diskursanalytisch orientierte Arbeiten bieten sich neben den Frequenzanalysen und den Analysen der Charakteristika eines Teilkorpus, auch die Analyse von Kookkurrenzen für eine an Laclau und Mouffe orientierte empirische Untersuchung an (Glasze 2013b: 106ff; Mattissek 2007: 89–92). Ähnlich wie die »Fixierung von Knotenpunkten« in der Theorie Laclaus und Mouffes stellen Lothar Lemnitzer und Heike Zinsmeister fest, dass der Kontext eines Textelements für die Generierung von Sinn und den Sprachgebrauch von zentraler Bedeutung sind und nicht abstrakte sprachliche Regeln (Lemnitzer/Zinsmeister 2006: 145ff; s. auch Dzudzek 2013: 79). Daran anschließend bietet sich deshalb für eine an Laclau und Mouffe anschließende Untersuchung auch »die Analyse empirisch vorliegender Kontexte, wie sie in Texten vorzufinden sind« an (Dzudzek 2013: 79). Wie Dzudek bezugnehmend auf Lemnitzer und Zinsmeister erklärt, wird unter »Kontext« die Summe der unmittelbaren Rahmenbedingungen einer Aussage als das Bezugssystem verstanden, »innerhalb dessen einer Äußerung eine Funktion zukommt« (ebd.: 79).

Zur empirischen Umsetzung der Untersuchung des Kontextes von Schlüsselbegriffen bieten sich *Kookkurrenzanalysen* an. Im Rahmen von diskursanalytischen Arbeiten, vor allem einer Arbeit wie sie hier in Anlehnung an Laclau und Mouffe ausgeführt wird, bilden sie einen weiteren nützlichen Schritt der Analyse. Mit Kookkurrenzanalysen ist es möglich, die Spezifität zu ermitteln, mit der bestimmte Textelemente miteinander verknüpft sind. Kookkurrenzen werden deshalb auch als »das gemeinsame Vorkommen zweier Wörter in einem gemeinsamen Kontext betrachtet« (Lemnitzer/Zinsmeister 2006: 147). Zur Umsetzung dieses Analyseschritts wird eine festgelegte Umgebung um bestimmte Schlüsselwörter herum extrahiert und nach dem Kriterium untersucht, welche Wörter mit einer überzufälligen Häufigkeit um das Schlüsselwort auftauchen. Die Umgebung des jeweiligen Schlüsselwortes ist durch eine vordefinierte Anzahl an Satzzeichen definiert, die im Korpus vor und nach dem entsprechenden Textelement verzeichnet sind (Dzudzek 2013: 79f). Diese Methode ermöglicht es herauszufinden, in welchem Teilkorpus, welche Begriffe in einen gemeinsamen Kontext gestellt werden. In Bezug auf die Theorie Laclaus und Mouffes lassen sich damit Aussagen über die »Differenz- und Äquivalenzbeziehungen« der lexikalischen Elemente zueinander treffen (ebd.: 80).

Bevor mit den Kookkurrenzanalysen begonnen werden kann, müssen sogenannte Schlüsselbegriffe freigelegt werden, um welche herum die Analyse von Kookkurrenzen angesetzt werden. Mithilfe der zuvor absolvierten CED werden die charakteristischen Begriffe des Diskurskorpus freigelegt, um welche stichprobenhaft Kookkurrenzanalysen durchgeführt werden, um Vermutungen über charakteristische Begriffe bestätigen zu können. In der hier vorliegenden Arbeit bezog sich die CED vor allem auf die Identifizierungen der einzelnen Sprecher*innenpositionen mit bestimmten charakteristischen Begriffen. Da die charakteristischen Begriffe in Anlehnung an Laclau und Mouffe sogenannte Knotenpunkte darstellen, sollen nun mithilfe der Kookkurrenzanalyse die signifikanten Äquivalenz- und Differenzbeziehungen zu anderen Textelementen bestimmt werden, um besser ihre »temporären Fixierungen« als Knotenpunkte des Diskurses verstehen zu können (Glasze 2013b: 82f).

Es handelt sich um signifikante Kookkurrenzen in Beziehung zu den Schlüsselbegriffen, wenn »deren Vorkommen deutlich die Wahrscheinlichkeit eines zufälligen Zusammentreffens übersteigen« (Scherer 2014: 48). Da es sich bei einem signifikanten gemeinsamen Auftreten i.d.R. um sinnvolle bedeutungsgenerierende Kookkurrenzen handelt, wird auch von Kollokationen gesprochen (ebd.: 48). Von einer äußerst deutlichen Wahrscheinlichkeit des Zusammentreffens kann gesprochen werden, wenn der in Kookkurrenz auftauchende Begriff, ebenso einen signifikanten CED-Wert aufweist.¹⁶ Je größer der CED-Wert, der in Kookkurrenz auftauchenden Worte, desto wahrscheinlicher ist es folglich, dass die beobachtete Häufung dieser Begriffe in den Aussagen des transnationalen Diskurses Globaler Bildung statistisch signifikant (d.h. überzufällig) ist (Dzudzek et al. 2009: 246).

Den Ausgangspunkt für die Umsetzung der Kookkurrenzanalysen stellen demnach die mit dem Programm Lexico gewonnen lexikometrischen Daten aus der CED dar.

16 Dies gilt insbesondere für die Kookkurrenz des charakteristischen Begriffs *citizenship* und seine antagonistischen Elemente *violence* und *discrimination* im Korpus der UNESCO (siehe Tab. 4).

Durch diesen Verfahrensschritt werden für jedes Teilkorpus, mögliche signifikante charakteristische Begriffe herausgearbeitet. Jedes Teilkorpus besteht in der hier vorliegenden Untersuchung aus je einer spezifischen Sprecher*innenposition, welche wiederum als modellhaftes Beispiel einer Sprecher*innenposition der GB in der Global Governance-Architektur ausgewählt wurde, womit deren Grad der Institutionalisierung auch im engen Zusammenhang mit der Materialität der sprachlichen Artikulationen steht.¹⁷ Zur Umsetzung der Kookkurrenzanalysen kann auch die Bildung von sogenannten Konkordanzen,¹⁸ um die erarbeiteten »Keywords«¹⁹ sehr hilfreich sein, da sie helfen auf Grundlage der besser ersichtlichen Kontextinformationen schneller musterhaftes gemeinsames Vorkommen erkennen zu können. Für jede Sprecher*innenposition des Diskurskorpus können so häufig wiederkehrende Kookkurrenzen, welche im Zusammenhang mit den charakteristischen Begriffen des Diskurses GB stehen, auf der Basis von Häufigkeiten quantitativ gestützt – anders gesagt: auf Basis der Materialität lexikalischer Zeichen – relativ zügig ermittelt werden. Die hier gewonnenen Erkenntnisse werden mit in die Feinanalyse am Ende der Untersuchung aufgenommen, da hier bereits Hinweise für die Qualität der Beziehungen der Schlüsselbegriffe mit anderen Elementen sichtbar werden.

-
- 17 Man kann in diesem Sinne auch von einer »mehrdimensionalen Materialität des Diskurses« sprechen, denn die »Materialität des Diskurses« besteht nicht nur in der Materialität der Zeichen, sondern auch in dem Grad der Institutionalisierung, mit dem die lexikalischen Zeichen in der Welt in den Umlauf gebracht werden. Wie bereits unter 3.2.1 zur »pädagogischen Globalisierung« erwähnt, gehören dazu eben auch die Macht über Ressourcen und Mittel über Publikationsstrukturen und -kanäle zu verfügen und Chancenmöglichkeiten in diesen zu veröffentlichen. Je höher der Grad der Institutionalisierung einer institutionalisierten Gruppe ist, desto höher ist zumeist auch die Verfügungsmacht über Ressourcen und Chancenermöglichungen.
- 18 Als Konkordanz werden sogenannte Listen bezeichnet, »die alle Vorkommen eines ausgewählten Wortes – oder auch Wortfolgen – in seinem Kontext« zeigen (Glasze 2013b: 107). Diese deshalb sogenannten *keywords in context* oder in Kurzform auch KWIC genannt, werden in der Regel in der Analyse zeilenweise dargestellt. Das Schlüsselwort, welches als Suchbegriff angegeben wurde wird dazu in der Mitte der Textzeile grafisch hervorgehoben. Links und rechts des Begriffs wird der ausgewählte Kontext der Zeile angegeben. Zudem lassen sich die davor und danach liegenden Konkordanzen auch in der alphabetischen Reihenfolge anordnen, welches die Analyse wesentlich vereinfacht (Lebart/Salem/Berry 1998: 32f; Scherer 2014: 44). Die Analysesoftware Lexico 5, aber auch MAXQDA ermöglichen diese Funktion. So kann mithilfe eines Computerprogramms auf der linken und rechten Seite um das ausgewählte *keyword* herum ein festgelegter Kontext bestehend aus einer vorher festgelegten Anzahl von Zeichen oder Wörtern angezeigt und alphabetisch geordnet werden. Konkordanzen eignen sich deshalb für die Vorbereitung der Kookkurrenzanalyse, aber auch zur Bereitstellung ausgewählter Textpassagen in der qualitativen Interpretation. Jeweils links und rechts der Schlüsselwörter erscheinen anschließend Tabellen mit an das Schlüsselwort angrenzenden Wörtern im Rahmen der jeweils zuvor festgelegten Zahl an Satzzeichen.
- 19 Scherer spricht bei Keywords auch von sogenannten »Knoten«, wodurch es mit der Theorie Laclaus und Mouffes, die von sogenannten Knotenpunkten sprechen, große Übereinstimmungen gibt (Scherer 2014: 44).

5.3 Analyse narrativer Muster als interpretative Feinanalyse ausgewählter Textpassagen eines offenen Diskurskorpus

5.3.1 Neue Ansätze der Erzählforschung als Analyseansatz

Quantifizierende Verfahren – wie lexikometrisch-korpuslinguistische Verfahren – reichen nach Glasze in der Regel vielfach nicht aus, um die Tiefe der Bedeutung von Texten nachvollziehen zu können, da mit lexikometrischen Verfahren zumeist nur die Verknüpfungen einzelner lexikalischer Elemente nachvollzogen werden können (Glasze 2013b: 113f). Wie er weiter ausführt, bedarf es dazu qualitativer Ansätze, die die vielfältigen Verbindungen und vielschichtigen Relationen von lexikalischen Elementen, welche oberhalb der Wort- und Satzebenen oder häufig sogar oberhalb der Ebene einzelner konkreter Texte auftreten, erfassen können. Für solch einen feinanalytischen Schritt im Rahmen von Diskursanalysen bietet sich dafür das stärker interpretative Kodieren von Elementen und Verknüpfungen an (ebd.: 114). Das Ziel des Kodierens ist es, die »Regelmäßigkeiten im (expliziten und impliziten) Auftreten (komplexer) Verknüpfungen von Elementen in Bedeutungssystemen herauszuarbeiten« (ebd.: 114). Wie Glasze hervorhebt, wird es damit möglich, sich Hinweise auf diskursive Regeln bewusst zu machen. Zur Umsetzung solcher Kodierungsverfahren kommen Techniken der qualitativen Sozialforschung, die im Zuge der Grounded Theory entwickelt wurden (Strauss/Corbin 1996), zur Anwendung. Diese müssen jedoch zuvor noch an die theoretischen Vorannahmen der Diskurstheorie nach Laclau und Mouffe angepasst werden. Konzeptionell kann das bewerkstelligt werden, indem im Zuge der kodierenden Verfahren auf die Überlegungen der neueren Erzähltheorie zurückgegriffen wird (Glasze 2013b: 114).

Mit dem Begriff »Erzähltheorie werden heterogene Ansätze der Erzählforschung beschrieben, die auf eine systematische Beschreibung der Formen, Strukturen und Funktionsweisen narrativer Phänomene abzielen.« (Nünning/Nünning 2002: 4) Ziel der Erzähltheorie ist es, eine systematische »Darstellung der wesentlichen Elemente des Erzählens und ihrer strukturellen Zusammenhänge« (Stanzel 1995 zit.n. Nünning/Nünning 2002: 4) aufzuzeigen. Ihre Grundlagen hat die Erzähltheorie im russischen Formalismus sowie dem Strukturalismus. Hier stand im Wesentlichen das Interesse an der erzählten Geschichte selbst im Vordergrund, womit das Bemühen einherging »eine möglichst abstrakte, eindeutige und kohärente Metasprache zur Beschreibung der Konstitution, Relationen und Strukturen narrativer Texte zu entwickeln.« (Nünning/Nünning 2002: 6) Mit der Zeit verlagerte sich das Interesse der Erzähltheorie durch die Interventionen Gérard Genette's (1980), jedoch immer mehr auf Fragen nach der Zeitstruktur und auf die Formen der erzählerischen Wiedergabe (Nünning/Nünning 2002: 7f).

Nünning verdeutlicht, dass die Erzähltheorie seit einigen Jahren nicht mehr die alleinige Domäne des Strukturalismus ist. Vielmehr kann man heute, wenn man über Erzähltheorie bzw. Narratologie spricht, diese Begriffe nur noch im Plural verwenden (Nünning/Nünning 2002: 9ff). So sind mit der Integration neuer Ansätze wie z.B. des Poststrukturalismus, aber auch durch neue Anwendungen und Weiterentwicklungen erzähltheoretischer Kategorien, Modelle, Methoden und Grundbegriffe, durch Frage- und Themenstellungen wie z.B. der *postcolonial*- und *gender*-Forschung, immer mehr die kon-

text- und themenbezogenen Zusammenhänge des Erzählens in den Mittelpunkt der Forschung gerückt (ebd.: 13f). Damit überschreiten sie »zum einen das Erkenntnisinteresse der dominant textzentrierten Narratologie strukturalistischer Provenienz« (ebd.: 13). Auf der anderen Seite sind durch die Untersuchung von Aspekten wie *ethnicity* (z.B. *postkoloniale Narratologie*), *class* (z.B. *marxistische Narratologie*), *gender* (z.B. *feministische Narratologie*) »zuvor vernachlässigte, thematische, ethische und ideologische Fragen in den Vordergrund« gestellt worden (ebd.: 13f). Wie Georg Glasze hervorhebt, ist die Öffnung der Erzählforschung parallel zu den Entwicklungen in den Sozialwissenschaften zu betrachten, da hier vermehrt die Konzepte der Erzähltheorie zur Untersuchung nicht-fiktionaler Texte Anwendung fanden (Glasze 2013b: 114).

Wurde zuvor die Erzählforschung größtenteils auf literarische Texte angewandt, entstanden ab den 1980er Jahren eine Vielzahl von Studien, die die Erzählforschung auf weitläufige Bereiche der Gesellschafts- und Sozialforschung²⁰ ausweiteten. Diese reichen von der Geschichts- über die Policy-, Ungleichheits- und Klimagerechtigkeitsforschung bis hin zu den Gender Studies und den Postkolonialen Studien (bspw. White 1980; Stone 1989; Somers 1994; Birk/Neumann 2002; Viehöver 2004, 2015).

Margarte R. Somers führt aus, dass die neuen Ansätze der Erzählforschung Narrationen als Konzepte sozialer Epistemologien und Ontologien²¹ begreifen: D.h., dass diese Konzepte beanspruchen, dass durch Narrationen erst unser Wissen und Sinnverständnis hervorgebracht wird, wodurch unsere soziale Identität und Wirklichkeit konstituiert wird. Egal wer wir sind – ob nun (Sozial-)Forschende oder die Erforschten selbst, Lernende oder Lehrende – alle von uns werden das sein, wozu wir durch Fremd- und Selbstverortungen in sozialen Narrationen, also jenen, die nur kaum von uns allein hervorgebracht wurden, geworden sind (Somers 1994: 606, 615). Die neuen Ansätze der Erzählforschung haben den Anspruch Identitätskonstruktionen im Zusammenhang mit »Artikulationen, Stimme und Handlungsermächtigung« in Abhängigkeit von Situation und Kontext zu analysieren (Birk/Neumann 2002: 121). Dies geht auf eine Weigerung der neuen Ansätze zurück, diese beiden Bereiche als getrennt zu betrachten, da hier ein Verständnis vorliegt, dass davon ausgeht, dass jede Identitätskonstruktion und soziale Wirklichkeit immer auch in unsere Artikulationen über unser Handeln und unsere Handlungsmacht eingeschrieben ist (Somers 1994, 614f) und deshalb die »Identitätskonstruktionen und -zuschreibungen die Grundlage für individuelle sowie kulturelle Selbstbestimmung oder aber Selbstentfremdung« bilden (Birk/Neumann 2002: 119). Soziale Identität ist damit als etwas zu begreifen, das niemals vollendet, vielmehr prozedural-dynamisch ist und deshalb kein Wesen, sondern nur temporäre Fixierungen innerhalb

20 Der Einzug der Erzählforschung in den Bereich der Sozialwissenschaften ist nicht ohne langwierige Debatten um die Legitimität dessen geschehen, was als Erkenntnis über Handlungsmöglichkeiten gilt. Lange Zeit galt in den Sozialwissenschaften die Ansicht, dass Handeln und Handlungsmacht nur über die Erforschung beobachtbarer sozialer Verhaltensweisen analysiert werden kann. Fragen der Identitätskonstruktion und Ontologie wurden deshalb anderen Disziplinen wie der Spekultativen Philosophie oder Psychologie überlassen (Somers 1994: 615, 1996).

21 Wichtig ist zu erwähnen, dass das nicht immer so war, sondern sich das Verständnis der Erzählforschung gewandelt hat. So gab es eine Verschiebung des Fokus der Erzählforschung von der repräsentativen Ebene hin zur ontologischen Ebene, welche sich dann erst auf die Sozialwissenschaften ausweiten konnte (Somers 1994: 613).

eines Zeitverlaufs hat (Birk/Neumann 2002: 121f). Worum es schließlich der neueren Erzählforschung geht, ist, die Analyse der von Zeit und Raum abhängigen Konstruktion von Identität und die Analyse von Handlungsmacht innerhalb eines zu untersuchenden Feldes zusammenzuführen (Somers 1994: 615; Birk/Neumann 2002: 121).

Es ist dabei jedoch hervorzuheben, dass es sich bei der Konstruktion von Identität immer um Konzepte einer relationalen Konstitution von Identität handelt, da hierbei stets von »einer notwendigen[,] aber letztlich immer unmöglichen Abgrenzung von »dem Anderen« ausgegangen wird« (Glasze 2013b: 115). So gehen die neuen Ansätze der Erzählforschung davon aus, dass »die Konstruktion des Eigenen [...] unauflöslich mit der Vorstellung des Anderen verknüpft ist« (Birk/Neumann 2002: 123). Ein Wir-Bewußtsein kann es deshalb nicht ohne irgendwelche anderen Gruppen oder das »Andere« geben (ebd.: 123). Die subjektive Idee vom Anderen, Andersartigen oder Fremden ist infolgedessen immer relational und darum »kann es bei der Frage nach dem Verhältnis zwischen dem Eigenem und Fremden nicht um das »Wesen« des Anderen gehen« (ebd.: 123f). Ergo stehen ebenso im Fokus der neuen Ansätze der Erzählforschung soziokulturell dominante Wahrnehmungsmuster und narrative Konstruktionen, die in den Diskursen zirkulieren, um die Konstruktion stereotypisierender Repräsentationen des Anderen beschreibbar zu machen und ihre Funktionen frei zulegen (ebd.: 123f). Auch das Verhältnis zwischen Eigenem und Fremdem ist stets dynamisch-prozessual und niemals vollendet (ebd.: 124). Gerade an diesem Punkt tritt der enge Bezug zur Diskurstheorie von Laclau und Mouffe hervor: Narrationen können demnach »als Artikulationen [...] analysiert werden, die eine Beziehung zwischen Elementen herstellen, Grenzen etablieren, auf diese Weise eine temporäre Fixierung leisten, Bedeutung und damit Identität konstituieren.« (Glasze 2013b: 115)

Texte stellen im Sinne der neueren Erzählforschung also keine Abbildung von Realität dar, sondern generieren als eine Art der Ausprägung kultureller Selbstverständigung soziale Wirklichkeitskonstruktionen. Sie sind als »Weisen der Welterzeugung« zu analysieren (Birk/Neumann 2002: 116). Identität wird damit als etwas gesehen, das sich im Wesentlichen über diskursive Formationen bildet und festigt – am Intensivsten über Narrationen (ebd.: 119f). Glasze schlägt deshalb vor, in Anlehnung an Somers, eine für die Sozialwissenschaften relevante Reinterpretation von Narrationen, welche für den feinanalytischen Schritt hegemonietheoretischer Diskursanalysen nach Laclau und Mouffe produktiv erscheinen, vorzunehmen (Glasze 2013b: 116). Somers weist darauf hin, dass Narrationen nur durch ihre netzwerkartigen Verknüpfungen, die sie in einem Beziehungsgeflecht aus symbolischen, institutionalisierten und materiellen Praktiken eingehen, verstanden werden können. Mithilfe des Begriffs *emplotment*, welchen Somers von Hayden White entlehnt,²² verdeutlicht sie, dass erst das Zusammenfügen von Elementen zu Bestandteilen einer versteh- und memorierbaren Narration, also durch eine narrative Geschlossenheit, eine Bedeutung der unabhängigen Elemente erreicht wird und nicht

22 Geschehnisse, Ereignisse und Handlungen für sich einzeln genommen sind tendenziell bedeutungsfrei, da sie für sich allein keine Erklärung des Geschehens anbieten. Für sich einzeln nebeneinander stehend sind sie eine »Chronik«. Erst durch das Zusammenfügen der Elemente der »Chronik« zu einer Geschichte (*story*) werden sie verstehbar, einsehbar, memorierbar; dieser Vorgang heißt bei White *emplotment*.« (Lexikon der Filmbegriffe 2012)

allein durch ihre chronologische oder kategorische Anordnung (Somers 1994: 616). Diese Verstehbarkeit im Rahmen einer Narration wird wiederum durch Beziehungen einer spezifischen Qualität zwischen den Elementen erreicht, wozu Beziehungen der Äquivalenz, der Opposition, der Kausalität oder der Temporalität zählen (Glasze 2013b: 116; Somers 1994: 616). Erst durch diese Beziehungen lassen sich dann auch Verbindungen zu anderen Narrationen herstellen. Die verschiedenen Formen der Beziehungen stellen somit eine Art Logik oder Syntax der einzelnen Narrationen dar (Somers 1994: 617). Es sind jene musterhaften Konfigurationen der qualitativen Beziehungen in den Narrationen, die schließlich Bedeutung konstituieren (ebd.: 617).

Somit lässt sich diskurstheoretisch im Anschluss an Laclau und Mouffe festhalten, dass narrative Muster »als regelmäßige Verknüpfungen von Elementen gefasst werden, die Beziehungen einer spezifischen Qualität herstellen« (Glasze 2013b: 115). Um die Hervorbringung von politischen Bedeutungen und Identitäten zu analysieren, stehen dabei die qualitativen Beziehungen der Äquivalenz, der Opposition, der Kausalität oder der Temporalität im Fokus hegemonietheoretischer Diskursanalysen (Glasze 2013b: 116–118; Somers 1994: 616). Wichtig ist zu beachten, dass Narrative Muster unablässig in umfassendere Narrationen eingebunden sind, welche in konkreten Texten nur teilweise reproduziert werden (Glasze 2013b: 115). Die Beziehung zwischen Narrationen und narrativen Mustern gestalten sich damit stets komplex und polyvalent, weil Bedeutungsfelder durch eine Vielzahl von narrativen Mustern vermittelt werden und Narrationen als Träger vieler aufeinander bezogener narrativer Muster fungieren (Birk/Neumann 2002: 116). In Rückbezug auf Edward Said unterstreichen Birk und Neumann die kontextbewusste Textanalyse postkolonial inspirierter Erzählforschung, weshalb es nach ihnen im Zuge von Analysen narrativer Muster unvermeidlich ist, »die kulturellen, sozialen wie historischen Entstehungsbedingungen« der verschiedenen Texte mit zu berücksichtigen (ebd.: 116).

Der hier vorliegende Analyseansatz der neueren Erzählforschung zur Analyse narrativer Muster ist eng mit dem Versuch der Offenlegung von Artikulations- und Handlungsmöglichkeiten verbunden, weil damit die bewusste Positionierung im gesellschaftlichen Kontext überhaupt erst ermöglicht wird. Erst durch das Bewusstwerden dieser Positionierungen im diskursiven Feld können daraufhin spezifische Formen der Artikulation gezielt als Medium im Sinne eines widerständigen Handelns eingesetzt werden (ebd.: 122f).

5.3.2 Kodierende Verfahren als Operationalisierung der Analyse narrativer Muster

Diskursanalysen, die an Foucault bzw. Laclau und Mouffe konzeptionell anschließen, haben zum Ziel die Konstitution und den Wandel von sprachlichen Bedeutungssystemen zu untersuchen. Es geht somit darum, die komplexen Bedeutungssysteme, die durch die vielschichtige Verknüpfung einzelner Textelemente hergestellt werden, zu untersuchen (Glasze/Husseini/Mose 2009: 293). Dazu müssen die »vielfältigen Verbindungen und vielschichtigen Relationen oberhalb der Wort- und Satzebene, häufig sogar oberhalb der Ebene einzelner konkreter Texte« in den Blick genommen werden (ebd.: 293). Als ein Teilschritt der Diskursanalyse dient dazu zur Umsetzung »das stärker interpretative Kodieren von Elementen und deren Verknüpfungen« (Glasze 2013b: 114) im An-

schluss an die konzeptionellen Überlegungen der neueren Erzähltheorie in Bezug auf die Analyse narrativer Muster. Um die Analyse narrativer Muster unter Hilfestellung der Analysekategorien postkolonialer Erzählforschung im Rahmen hegemonietheoretischer Diskursforschung nach Laclau und Mouffe operationalisieren zu können, bietet es sich an, die entsprechend angepasste Verwendung kodierender Verfahren nach der Methode des *theoretical sampling* der *Grounded Theory*, nach Strauss und Corbin (1996) vorzunehmen (Glasze 2013b: 115–117).

5.3.2.1 Kodieren in Diskursanalysen

»Kodieren« bezeichnet allgemein, wie es Strauss und Corbin definieren, all jene Vorgehensweisen »durch die die Daten aufgebrochen, konzeptualisiert und auf neue Art zusammengesetzt werden.« (Strauss/Corbin 1996: 39) In der Forschungspraxis handelt es sich dabei um das Markieren von Textstellen mit Begriffen, sogenannten »Codes«, die »entweder aus den Daten selbst stammen oder die der Forscher »von außen«, aus seinem (theoretischen) Vor- und Kontextwissen an die Daten anlegt.« (Diaz-Bone/Schneider 2004: 465f) Dahinter steht die Idee, mithilfe von empirischen Indikatoren, die im Analyseprozess durch die durch Begriffe markierten Textstellen präsentiert werden, »Möglichkeiten für eine Konzeptionalisierung, ein Fassen dieser Indikatoren in sie bezeichnende Begriffe (concepts)²³ zu identifizieren.« (ebd.: 466) Der Analyseprozess²⁴ verläuft dabei, indem eine Vielzahl solcher Indikatoren durch die forschende Person »begrifflich er- bzw. bearbeitet und systematisiert (codiert) – d.h. dimensionalisiert, klassifiziert, gruppiert usw.« werden. Im Zentrum des Prozesses steht ein ständiges Vergleichen und Abgleichen der begrifflich-theoretischen Arbeit an den Daten oder anders gesagt »an der begrifflichen Ordnung des untersuchten Feldes.« (ebd.: 466) Erreicht werden soll eine Verdichtung der in den Daten identifizierten Indikatoren

23 Konzepte stellen bei Strauss und Corbin (1996: 45) die grundlegenden Analyseeinheiten des Kodierungsprozesses dar. Den Prozess zur Herstellung der Konzepte, bezeichnen sie als Konzeptualisierung. In der an die *Grounded Theory* angelehnten Methode stellt das Konzeptualisieren der Daten den ersten Schritt innerhalb des Vorgangs der qualitativen Analyse dar (Strauss/Corbin 1996: 45). Demnach bezeichnet das Aufbrechen und Konzeptualisieren »das Herausgreifen einer Beobachtung, eines Satzes, eines Abschnitts und das Vergeben von Namen für jeden einzelnen darin enthaltenen Vorfall, jede Idee oder jedes Ereignis – für etwas, das für ein Phänomen steht oder es repräsentiert.« (ebd.: 45) Wichtig ist, zu erkennen, dass es sich beim Konzeptualisieren nicht um eine Zusammenfassung der Daten handelt, sondern um die Benennung der darin auftauchenden Phänomene (ebd.: 46).

24 Für den kodierenden Analyseprozess stellen Strauss und Corbin den Analysierenden mehrere Techniken bereit, die ein sehr sorgfältiges Kodieren der Daten ermöglichen. Grundlegend haben sie dazu, das Kodieren in drei Haupttypen aufgeteilt: Das a) offene Kodieren, b) axiale Kodieren und c) selektive Kodieren (Strauss/Corbin 1996: 40). Wie Strauss und Corbin (1996: 41) selbst feststellen, implizieren die Techniken nicht, dass an ihnen rigide festgehalten werden muss. Vielmehr geht es darum, diese Techniken »flexibel, den Umständen entsprechend anzuwenden.« (ebd.: 41) Da für die Kodierung die *Grounded Theory* nur als Orientierung für den Analyseprozess gilt und dieser auch aufgrund unterschiedlicher theoretischer Vorannahmen wesentlich modifiziert wurde, wird nachfolgend von diskursanalytischen Kodieren gesprochen und nicht von *Grounded Theory* selbst.

zu abstrakteren Konzepten und der Offenlegung ihrer »generalisierten empirischen Merkmale und handlungsrelevanten Deutungsbezüge untereinander« (ebd.: 466).

Um von einem diskursanalytischen Kodieren zu sprechen, bedarf es nach Rainer Diaz-Bone und Werner Schneider (2004: 466) einer diskurstheoretisch begründeten »Ausbuchstabierung« dessen, worauf sich das Kodieren im Rahmen der eigenen diskurstheoretischen Prämissen bezieht und woraufes zielen soll (ebd.: 466). Eine mögliche Form für die hier vorliegende Analyse bieten Glasze et al. an: »Das Kodieren von Textelementen und deren Verknüpfungen« dient in diskurstheoretisch orientierten Arbeiten nach Laclau und Mouffe dazu, »die Regelmäßigkeiten im (expliziten und impliziten) Auftreten (komplexer) Verknüpfungen von Elementen in Bedeutungssystemen herauszuarbeiten«, um von diesen Regelmäßigkeiten, Schlussfolgerungen auf die Regeln der diskursiven Bedeutungskonstitution ziehen zu können (Glasze/Husseini/Mose 2009: 293; s. auch Glasze 2013b: 116). Kodierende Verfahren stellen, diskursanalytisch angewendet, eine Technik zur Aufdeckung von Mechanismen zur Herstellung sozialer Wirklichkeit dar (Glasze/Husseini/Mose 2009: 294).

Zunächst wurden in den Sozialwissenschaften Kodierungstechniken, wie Glasze, Husseini und Mose betonen, »vor dem Hintergrund interpretativ-hermeneutischer Theorien entwickelt, bspw. im Rahmen der »qualitativen Inhaltsanalyse« (Mayring 2008) und Ansätzen der *Grounded Theory* (Glaser/Strauss 1998; Strauss/Corbin 1996)« (Glasze/Husseini/Mose 2009: 294). Was diese Ansätze der interpretativ-hermeneutischen Tradition eint, ist, dass das Kodieren hier dazu dient, Textstellen zu klassifizieren und zu bündeln, um dann mit der Vergabe der dabei entwickelten Codes Aussagen über den Sinn des Inhalts der untersuchten Textstellen zu tätigen. Codes stehen hier als Indikatoren für einen bestimmten Inhalt der Textstellen, mithilfe derer eine Interpretation über ihren Sinn geleistet wird (ebd.: 294).

Am wohl eindeutigsten lässt sich gemäß Glasze die Problematik einer solchen Herangehensweise am Kommunikationsmodell »Sender $\Rightarrow$ Inhalt $\Rightarrow$ Empfänger« der Qualitativen Inhaltsanalyse zeigen. Hier wird davon ausgegangen, dass vom »Inhalt« Rückschlüsse auf die »soziale Wirklichkeit« gezogen werden können – eine »soziale Wirklichkeit« die in Form des »Kommunikator[s]«, des »Rezipienten« oder der »Situation« auftritt. Zur Umsetzung solcher Inhaltsanalysen wird als wichtigste Methodik die Kodierung der »Inhalte« von Texten mittels eines Categoriesystems eingesetzt (Glasze 2013b: 99). Die Annahme die hier der qualitativen Inhaltsanalyse zugrunde liegt besteht darin, »dass jeder Text(-teil) einen ›Inhalt‹, d.h. eine Bedeutung transportiert und das dieser ›Inhalt‹ durch den Inhaltsanalytiker erschlossen werden kann.« (ebd.: 99) Aus diskursanalytischer Perspektive poststrukturalistischer Prägung ist eine solche Herangehensweise in zweierlei Hinsicht problematisch. Zum einen wird das hier verwendete Repräsentationsmodell, mit dem davon ausgegangen wird, dass es analytisch möglich wäre, Rückschlüsse von einem Text(-teil) auf die Bedeutung des Textes zu ziehen, grundlegend in Frage gestellt (ebd.: 99). Zum anderen steht auch die Arbeit mit Categoriesystemen in der Kritik, da diese das Risiko mit sich bringen »Tautologien zu erzeugen, indem ein etabliertes System durch Belegstellen reifiziert wird.« (ebd.: 99)

Den oben beschriebenen Unterschieden zufolge bedarf es einer von den hermeneutisch-interpretativen Ansätzen abweichende Gebrauchsweise der Kodierung im Rahmen

diskursanalytischer Arbeiten. So dienen in diskurstheoretisch geprägten Analysen die »Codes«

»zwar ebenso der Identifizierung (»Markierung«) von Textstellen, stellen jedoch nicht [...] als gehaltvolle Indikatoren (im Sinne des Konzept-Indikator-Verhältnisses) den notwendigen (!) Weg hin zu Schlüsselkonzepten dar [...]. Vielmehr müssen die jeweiligen »Codes« (als Verweise auf die in den Daten materialisierte Diskursordnung als »Realität sui generis«) entsprechend ihrer empirisch rekonstruierbaren »Verwendungsweisen« zu empirisch begründeten, diskurstheoretischen Aussagen über die Strukturiertheit, [die] Regelmäßigkeit dieser Ordnung zusammengefügt werden.« (Diaz-Bone/Schneider 2004: 474)

Anders gesagt obliegt den »Codes« in diskurstheoretisch geprägten Arbeiten nicht die Funktion, aus den »Codes« am Ende höhere Konzepte zu entwickeln, die einen tieferen Einblick in die einzelnen Aussagen geben können. Die »Codes« hier sind schlicht Markierungen, die Auskunft über die Einzigartigkeit der Charakteristika der spezifischen Aussage geben. Das Kodieren dient demnach eher als ein Verweis mithilfe dessen die Regelmäßigkeiten »anhand der [in] Textstellen sich materialisierenden diskursiven Praxis (typische Verwendungsweise, erkennbar aktualisierte Opposition u.ä.) [...] schnell und bequem herausgearbeitet werden [können].« (ebd.: 470) Die bestehenden Regelmäßigkeiten und Verbindungen zwischen Diskurselementen können somit schrittweise durch die zunächst entwickelten Kodierungen miteinander in Beziehung gesetzt werden. Die weiteren »Codes«, mit denen ein »Code« vernetzt ist, geben Einblick, welche »Wissenskonzepte« mit welchen Themen und Problematisierungen »vernetzt« sind, um so »thematische Regionen« herauszuarbeiten (ebd.: 470f). Der Unterschied zwischen der Kodierung in hermeneutisch-interpretativen Ansätzen und denen innerhalb diskurstheoretisch geprägter Analysen liegt, gemäß Glasze, demnach wesentlich im konzeptionellen Stellenwert des Kodierens und weniger im Ablauf der Kodierung (Markierung, Ordnung, Klassifizierung), da letzteres bei beiden vielfach ähnlich verläuft (Glasze 2013b: 116). Kurz gesagt kommt dem Kodieren innerhalb diskurstheoretischer Arbeiten ein schwächerer Status zu, wie Diaz-Bone und Schneider (2004: 474) hervorheben.

5.3.2.2 Gegenstand der Kodierung in Diskursanalysen nach Laclau und Mouffe

Bis hierher ist deutlich geworden, dass es das Ziel des Kodierens ist, die »Regelmäßigkeiten in den Beziehungen zwischen lexikalischen Elementen bzw. Konzepten in Diskursen herauszuarbeiten, um damit [...] Regeln der Konstitution von Bedeutung« in dem zu untersuchenden Feld ableiten zu können (Glasze 2013b: 116). Nun ist es aber von Wichtigkeit, wie Glasze betont, die daran anschließende Frage, was eigentlich kodiert werden kann, in zwei Schritte zu unterteilen. So unterscheidet er zwischen »dem Kodieren selbst« und »der Analyse der Regelmäßigkeiten, die sich im Überblick über die kodierten Textstellen erkennen lassen.« (ebd.: 116) Sein Vorschlag für ein gelingendes Kodieren ist es dabei konzeptionell zwischen zwei Ebenen des Diskursiven zu unterscheiden: Den Elementen und den Artikulationen.

Elemente

Die Basiseinheit eines Diskurses bilden nach Laclau und Mouffe sogenannte Elemente. In der Forschungspraxis müssen diese, wie George Glasze, Shadia Hussein und Jörg Mose hervorheben, jedoch in zweierlei Hinsicht unterschieden werden. Einerseits können Elemente als lexikalische Elemente im Sinne von einzelnen Wörtern bzw. Wortfolgen – wie z.B. *citizenship* oder *sustainable development* – verstanden werden. Andererseits ist es auch möglich diese als semantische Konzepte aufzufassen (Glasze/Hussein/Mose 2009: 295). Dies ist immer dann der Fall, wenn durch mehrere Signifikanten als Bedeutungsträger, also mehrere Bezeichner, ein einzelnes gemeinsames Signifikat als Inhalt, ein Bezeichnetes, bezeichnet werden kann. Ein Beispiel für solch ein semantisches Konzept ist *Bildung für Nachhaltige Entwicklung* und *Education for Sustainable Development*, da beide das gleiche semantische Konzept transportieren, obwohl lexikalisch unterschieden.

Im Zuge des hier angewandten Kodierungsprozesses dienen Elemente als Suchrasen für die Kodierung: D.h. sie selbst stellen keine »Codes« dar, sind jedoch Bestandteil eines Codes. Ihr Auftauchen wurde durch das zuvor angelegte und oben beschriebene lexikometrische Verfahren bereits ermittelt (ebd.: 295). Die Suche nach Elementen bezieht sich in der hier an Laclau und Mouffe angelehnten Diskursanalyse damit zunächst im Wesentlichen auf die charakteristischen Begriffe des zu untersuchenden diskursiven Feldes. Dies ist damit begründet, dass durch ihre analytisch festgestellte Charakteristik davon auszugehen ist, dass diese im zu untersuchenden Feld, durch die Gebrauchsweisen der Sprecher*innenpositionen hegemonial geworden sind. Die als charakteristisch klassifizierten Elemente werden nach ihrer Ermittlung in den Texten des Diskurses aufgesucht. Der Kontext dieser durch die lexikometrische Analyse erarbeiteten Signifikanten dient somit als Ansatzpunkt für die Analyse der narrativen Muster, da aus arbeitsökonomischen Gründen nicht die Gesamtheit des Textkorpus untersucht werden kann (Glasze 2013b: 132).

Artikulationen

Wie oben bereits erwähnt, zielt das Kodieren nicht auf die einzelnen Elemente ab, sondern auf ihre Verknüpfungen untereinander. Glasze (2013b: 116) schlägt zur Beschreibung der Verknüpfungen das »Konzept der Artikulation« vor, einen Begriff, der der Theorie Laclaus und Mouffes entnommen ist. In Anlehnung an Somers (1994: 616) setzen nach Glasze »Artikulationen Elemente miteinander in Beziehung und stellen auf diese Weise Beziehungen einer spezifischen Qualität her – bspw. Beziehungen der Äquivalenz, der Opposition, der Kausalität oder der Temporalität.« (Glasze 2013b: 116) So können allein schon zwei miteinander verbundene Elemente eine Verknüpfung herstellen, jedoch werden zumeist »komplexe Verbindungen zwischen verschiedenen Elementen gebildet.« (ebd.: 116) Artikulationen stellen demnach Verbindungen bzw. Verknüpfungen dar. Glasze bezeichnet diese komplexen Verknüpfungen als »narrative Muster«, weist jedoch im selben Moment darauf hin, dass sie vielerlei Ähnlichkeiten zu den andernorts als *plots* (Viehöver 2011: 215ff) oder *storyline* (Hajer/Waagenar 2003) bezeichneten Muster- oder Strukturprinzipien von Narrationen haben (Glasze 2013b: 116). Die Untersuchung von narrativen Mustern hat also zum Ziel, die Verknüpfungen der Elemente der narrativen Muster aufzudecken, damit aus den Regelmäßigkeiten der Beziehungen der kom-

plexen Verknüpfungen zwischen den verschiedenen Elementen die Regeln der Konstitution von Bedeutungen sichtbar werden. Denn es sind die narrativen Muster und »nicht das chronologische Auftreten von Ereignissen, Akteuren, Objekten etc., [...] die der Geschichte Sinn, Kohärenz, zeitliche und räumliche Strukturen verleih[en] und Beziehungen zwischen Ereignissen, Akteuren und Orientierungen herstell[en]«, um so schließlich Bedeutungen zu konstituieren (Viehöver 2011: 203). Narrative Muster müssen demnach als »central organising principle« innerhalb von Narrationen gesehen werden (ebd.: 203).

5.3.2.3 Diskursanalytisches Kodieren zur Freilegung narrativer Muster

Bei der Kodierung narrativer Muster kann grundlegend induktiv sowie deduktiv vorgegangen werden. Bei einer induktiven Vorgehensweise wird das Kategorie- oder Codesystem beim Durchgang durch das Material im Laufe des Kodierens selbst entwickelt. In einem offenen Prozess werden die einzelnen Codes während des Kodierens den dabei gewonnen Erkenntnissen nach modifiziert und angepasst (Glasze/Husseini/Mose 2009: 296). Demgegenüber stehen deduktive Herangehensweisen. Sie greifen auf ein Kategorie- oder Codesystem zurück, welches bereits vor der Arbeit am Textmaterial zur Verfügung steht. Auf Grundlage theoretischer Annahmen können so bestimmte Artikulationen als Code markiert werden, um daraufhin gezielt im Textkorpus danach zu suchen (ebd.: 296f). Darüber hinaus gibt es jedoch noch eine weitere Möglichkeit zur Bestimmung von Codes. Dieser dritte als diskursanalytischer Kodierungsprozess bezeichnete ist allgemein in drei Teilschritte zu unterteilen, welche sich an den Ausführungen von Diaz-Bone und Schneider orientieren (Diaz-Bone/Schneider 2004: 466):

1. Die ersten Codes verweisen auf die relevanten Grundelemente der diskursiven Praktiken in den entsprechenden Textpassagen im zu untersuchenden Feld. Damit sind zunächst Begriffe und Wortfolgen, aber auch Objekte, Sprecher*innenpositionen und Strategien gemeint, die in den Texten zu finden sind.

2. Im zweiten Schritt der Kodierung werden bereits Regelmäßigkeiten der qualitativen Beziehungen der verschiedenen Elemente herausgearbeitet. Im Rückgriff auf die Diskurstheorie Laclaus und Mouffes wird dabei, vor allem nach Beziehungen der Äquivalenz, der Opposition, der Kausalität oder der Temporalität gesucht, damit diese auf »manifeste oder latente Ebene benannt und rekonstruiert werden können« (ebd.: 466).

3. Im letzten Schritt der Kodierung stehen entlang der identifizierten und codierten Diskurselemente sowie den kodierten Regeln entsprechend, ihre Verknüpfungen untereinander im Fokus (ebd.: 466). Es geht darum die übergeordneten narrativen Muster durch den Kodierungsprozess erkennbar zu machen. Hier kann dann nach typischen Äquivalenz- und Oppositionskonstitutionen der Narrationen gesucht und ihre Muster können dann anhand von Codes herausgearbeitet werden.

Nachfolgend werden die einzelnen Teilschritte der hegemonietheoretischen Diskursanalyse noch einmal ausführlicher erörtert.

1. Kodier-Schritt: Auswahl der zu kodierenden Textstellen und die Erstellung eines ersten Codebuchs
Im ersten Schritt in der hier vorliegenden Arbeit, welche sich an die Operationalisierung der Diskurstheorie Laclaus und Mouffes nach Glasze anlehnt (Glasze 2013b), werden zunächst die Ergebnisse der lexikometrischen Analysen für die Definition von Codes

herangezogen. Auf der Ebene der Sprachoberfläche werden mithilfe der CED, charakteristische Begriffe des Diskurses sowie den Verlauf des Diskurses bestimmende Teilkorpora erarbeitet. Die Analyse der Kookkurrenzen spezifischer Wörter und Wortfolgen mit den charakteristischen Begriffen kann neue Codes freilegen (Glasze/Husseini/Mose 2009: 297). Das daran anschließende Verfahren des Kodierens dient der Klärung, welche Qualitäten die Verbindungen der lexikometrisch ermittelten charakteristischen Wörter bzw. Wortfolgen mit anderen Elementen aufweisen (ebd.: 297), also inwiefern »sie als Knotenpunkte dienen, welche Äquivalenzbeziehungen sie herstellen, ein Außen definieren und auf diese Weise Gemeinschaft konstituieren.« (Glasze 2013b: 117) Anders gesagt dient der Schritt des Kodierens des Kontextes der lexikometrisch ermittelten Begriffe dazu, Hinweise auf *zentrale Deutungsrahmen* und *Elemente der Narrationen* zu finden.

Dieser erste Schritt des Kodierens findet demnach hauptsächlich an der Diskursoberfläche an den sogenannten *Elementen* der ausgewählten Textpassagen statt. Hierunter sind »Begriffe«, »Objekte«, »Sprecher*innenpositionen«, »Konzepte«, »Ereignisse« und »Strategien« zu verstehen (Diaz-Bone/Schneider 2004: 466). Dies dient in erster Linie dazu das sogenannte *Thema* des Textes des Diskurses, also die Fakten, Begriffe, Ereignisse, Objekte und Konzepte sichtbar zu machen, um die der Text im Wesentlichen organisiert ist. Außerdem kann so die Problemstruktur herausgearbeitet werden, die der themenspezifische Diskurs konstituiert. Willy Viehöver schlägt deshalb in Anlehnung an die *Grounded Theory* das Stellen von Leitfragen vor, die in diesem Schritt an den Text gestellt werden können (Viehöver 2011: 208). Ihre Beantwortung, wie Viehöver betont, unterstützt die Bildung und Erstellung eines »Codebuchs« als ein erstes Aufbrechen der Textpassagen in ihre Elemente. Das Setzen von *Markern* (*Codes*) in den Textpassagen um die charakteristischen Begriffe herum, kann durch die Beantwortung z.B. folgender Fragen gefördert werden:

- Was geschieht im Text?
- Auf welche Begriffe, Objekte, Ereignisse, Sprecher*innenpositionen etc. weist eine Textpassage hin?
- Welche Phänomene werden wieder und wieder in den Daten wiedergegeben?
- Was ist das Thema und was ist das Hauptproblem? Wer definiert es wie?
- Welche Strategien werden zur Lösung im Text erwähnt?

2. Kodier-Schritt: Kodieren der Beziehungen der einzelnen Elemente und ihrer Qualität

Ist der erste Schritt, das Kodieren der wesentlichen Elemente, die die Textpassagen organisieren, geschehen, kann sich nun der Kodierung der Beziehungen der einzelnen Elemente in ihrer Artikulation untereinander gewidmet werden. Hierbei geht es vor allem darum, die verschiedenen Qualitäten der Beziehungen der Elemente untereinander aufzudecken. Bei der Kodierung zur Analyse narrativer Muster nach Laclau und Mouffe sind es dabei vor allem die qualitativen Beziehungen der Äquivalenz, Opposition, Kausalität und Temporalität (Glasze 2013b: 116), die im Mittelpunkt der weiterführenden Code-Bildung stehen. Der Schwerpunkt der Untersuchung liegt damit letztendlich auf den Fragen nach der Konstitution von Identität. Identität wird dabei immer als relationale Konstitution verstanden, weil dabei stets »von einer notwendigen, aber letztlich unmögli-

chen Abgrenzung von ›dem Anderen‹ ausgegangen wird.« (Glasze 2008: 205) Die Untersuchung von Narrationen kann im Sinne der Diskurstheorie nach Laclau und Mouffe mithilfe des Konzeptes der Artikulation durchgeführt werden, da durch sie die Herstellung von Äquivalenzen zwischen den Elementen, die Etablierung von Differenzen und Grenzen zwischen ihnen, ihre temporäre Fixierung sowie daraus entstehende Bedeutungen dieser analysiert werden kann. In seiner Gesamtheit gibt das Erfassen der qualitativen Beziehungen schließlich Aufschluss über die Frage der Konstitution von Identität innerhalb des themenspezifischen Diskurses (ebd.: 205).

In praktischer Hinsicht wird sich zur Erfassung der Verknüpfungen der Elemente in diesem Schritt ebenso der Erstellung von Codes bedient. Diese zielen jedoch diesmal darauf, die Beziehungen zwischen den zuvor identifizierten Elementen in den entsprechenden Textpassagen genauer erfassen zu können. Auf der manifesten oder latenten Ebene des Textes soll hier anhand der qualitativen Verweise der Elemente aufeinander die Art und Weise ihrer Beziehungen herausgestellt werden, um diese später in einem umfassenden Maße rekonstruieren zu können (Diaz-Bone/Schneider 2004: 466;480). Mithilfe der nun stattfindenden Kodierung der Merkmale der Beziehung zwischen bestimmten Wörtern bzw. Wortfolgen (den Elementen) durch die Erforschung ihrer grammatikalischen Konstruktionen zueinander soll nun bereits auf der Sprachoberfläche nach Hinweisen auf die Qualität der für Diskursanalysen nach Laclau und Mouffe wichtigen Beziehungen zwischen den Elementen (Äquivalenz, Differenz, Kausalität und Temporalität) gesucht werden. In Tabelle 1 (Tab. 1) werden im Anschluss an Glasze et al. Beispiele für grammatikalische Konstruktionen genannt, die Hinweise auf eine spezifische Qualität der Beziehung liefern können. Letztendlich bleibt jedoch die Feststellung der Qualität der Verbindungen zwischen den Elementen ein hochgradig interpretativer Schritt (Glasze/Husseini/Mose 2009: 297).

Tabelle 1: Hinweise für Beziehungen einer spezifischen Qualität zwischen Elementen.

Äquivalenz	Opposition (Antagonistisch)	Kausalität	Temporalität
»damit«, »für«, »ist«, »beitragen«, »fördern«, »darin bestehen«, »untrennbar mit etwas verbunden«, »um etwas zu tun können« stellen mögliche Hinweise dar für Äquivalenzbeziehungen	»entgegenstellen«, »Nein zu...«, »Die Bedrohungen durch...« als Hinweise für Oppositionsbeziehungen	»weil« oder »infolge« sind Hinweise für Kausalverknüpfungen	»danach«, »Nachdem«, »Bevor« sind Hinweise auf temporale Verknüpfungen

Quelle: Mit kleinen Veränderungen übernommen von Glasze/Husseini/Mose 2009: 297.

Die Spezifikation der qualitativen Beziehungen der Verknüpfungen lässt sich auch in diesem Schritt mithilfe weiterer Fragen und narratologischer Kategorien an den entsprechenden Text vereinfachen. Vor allem, wenn die Fragen auch auf die kontextuellen Bedingungen und Konsequenzen des Themas und des Hauptproblems in den Texten²⁵ abzielen (Birk/Neumann 2002: 129f; Diaz-Bone/Schneider 2004: 466f; Viehöver 2011: 208f). Mögliche Fragen an den Text, welche sich an den Ausführungen Diaz-Bones und Schneiders (2004: 466) und Viehövers (2011: 209f) sowie Birks und Neumanns (2002: 130ff) orientieren, können diesen Kodierschritt unterstützen:

- Welche Begriffe konstituieren welche Sprecher*innenpositionen?
- Welche Objekte und Ereignisse werden durch welche Begriffe und Strategien der Bezeichnung konstituiert?
- Wie werden Begriffs-, Objekt- und Ereignisgruppen in den diskursiven Verknüpfungen erstellt?
- Welche Sprecher*innenstrategien und -konzepte finden sich und welche Prozeduren müssen Sprecher*innen einsetzen, um ihren Beiträgen Gewicht und Geltung zu kommen zu lassen?
- Welche typischen Handlungsrezepte und Lösungen werden auf solchen Begriffs-, Objekt- und Ereignisgruppierungen von welchen Sprechern*innen errichtet?
- Zu welchen zeitlichen Abschnitten oder Ereignissen werden welche Objekte durch welche Begriffe und Bezeichnungen konstituiert?
- Welche Gemeinsamkeiten werden durch die Begriffe und Strategien der Sprecher*innen zu welchen Objekten hergestellt? Wie lassen sich die Gemeinsamkeiten identifizieren?

25 Um die Tiefendimensionen bei der Herausarbeitung der benannten Beziehungen spezifischer Qualität im Diskurs zu erhöhen, werden ebenfalls einige Untersuchungskategorien der postkolonialen Erzählforschung zur Generierung der Kodierung zur Unterstützung herangezogen (Birk/Neumann 2002: 130–147). Dies ergibt sich daraus, dass die postkoloniale Erzähltheorie besonders geeignete narratologische Kategorien bereitstellt, um die narrative Vermittlung von Identitätsentwürfen (ebd.: 130) im Spannungsfeld zwischen Äquivalenz und Opposition in einer spezifischen Raum-Zeit-Struktur zu analysieren. Dies geschieht vor allem in Rückbindung an die Fragestellung der hier vorliegenden Arbeit, da diese den Anspruch erhebt, der diskursiven Konstitution des in den Global Governance Strukturen generierten Diskurses GB mit einer höchstmöglichen Wachsamkeit gegenüber den postkolonialen Verstrickungen analytisch zu begegnen. Narratologische Kategorien, die das kritische Potenzial der Befragung des Textes stärken könnten, sind unter anderem *narration*, *focalization*, *Figurendarstellung* (z.B. der Lehrenden und Lernenden bzw. der Zielgruppen der Bildungsbestrebungen), *Perspektivenstruktur* (der Sprecher*innenpositionen), *Raumdarstellungen*, *Zeitstruktur*, *intertextuelle Referenzen* sowie die *qualitative Analyse der Sprache der einzelnen Ebenen der Sprecher*innenposition* entlang von zentralen Kategorien der Verortung von Identität wie Ethnizität, Gender, Klasse, Nation (ebd.: 130–147). Die Untersuchungskategorie der Figurendarstellung wird hier als Figurendarstellung bezeichnet, da dies eine Kategorie ist, die zunächst innerhalb der Literaturwissenschaft entwickelt wurde, um damit die im Roman vorkommenden Figuren analysieren zu können. Im übertragenen Sinne können im erziehungswissenschaftlichen Bereich damit die in den Texten beschriebenen Subjektgruppen, wie z.B. der Lehrenden und Lernenden beschrieben werden.

- Welche Gegensätze werden durch die Begriffe und Strategien der Sprecher*innen zu welchen Objekten hergestellt? Wie lassen sich die Gegensätze identifizieren?
- In welcher Beziehung stehen die einzelnen Sprecher*innen zueinander?
- Welche Raum-Zeit-Strukturen und Kausalitäten entstehen dabei?
- Geschehen die Artikulationen in einem selbstbewusst in den Diskurs eingreifendem Umfeld oder sind die Möglichkeiten begrenzt?
- Wie ist die Wahrnehmung über die gesprochen wird gestaltet? Werden diese konzeptionalisiert?
- Wie offen oder geschlossen sind die Perspektiven der Sprecher*innenpositionen?
- Welche Raumdarstellungen werden in den Artikulationen der Sprecher*innen wirksam?
- Welche Organisationen bzw. Anordnungen von Elementen der Artikulationen kehren stets wieder?
- Welche intertextuellen Referenzen tauchen immer wieder über den einzelnen Text hinaus auf? Inwiefern fungieren diese als Prätexte?
- Welche Rolle spielen die soziokulturellen Kategorien Ethnizität, Klasse, Gender, Nation, Globaler Norden oder Globaler Süden der Sprecher*innenpositionen? Welchen soziokulturellen Figuren wird in den Texten versucht eine Stimme zu verleihen?
- Welches Wirkungspotenzial haben diese Prätexte? Selbstbewußt hegemonial oder begrenzt?
- Welche qualitativen Merkmale weist die Sprache der Sprecher*innenpositionen und die Sprache der darin auftauchenden Figuren in Bezug auf die zentralen Kategorien der Identitätskonstruktion sowie der zentralen soziokulturellen Kategorien Ethnizität, Gender, Klasse, Nation, Globaler Norden, Globaler Süden auf?
- Gibt es Strategien der Subversion und Enthierarchisierung in der Sprache der Sprecher*innenpositionen? Wenn ja, welche?

3. Kodier-Schritt: Herausarbeiten der Regelmäßigkeiten mithilfe von Häufigkeitsanalysen und die Ableitung diskursiver Regeln

Nach dem Schritt des eigentlichen Kodierens müssen nun die Regelmäßigkeiten der Beziehungen innerhalb der kodierten Textstellen des Materials analysiert und herausgearbeitet werden. Dazu lassen sich die vorher angefertigten Kodierungen verwenden, um die Regelmäßigkeiten nach den expliziten Verknüpfungen von Elementen zu ordnen. Zur Analyse der Regelmäßigkeiten sind zunächst einmal Häufigkeitsanalysen der bereits erstellten Codes über die jeweiligen Verknüpfungen sehr hilfreich, um festzustellen wie häufig diese im Material vorkommen (Glasz/Husseini/Mose 2009: 298). Wie Glasze et al. hervorheben, können so »Rückschlüsse auf die Dominanz oder Marginalität bestimmter expliziter Verknüpfung von Elementen« gezogen werden.« (ebd.: 298) Häufigkeitsanalysen der Codes lassen sich im Anschluss als synchroner Vergleich analysieren, indem die Regelmäßigkeiten unterschiedlicher Korpora miteinander verglichen werden.²⁶

26 Andererseits wären aber auch diachrone Vergleiche möglich. Zur Realisierung diachroner Verfahrensweisen müssten die Regelmäßigkeiten entlang ihrer zeitlichen Verläufe genauer analysiert werden, indem mit Häufigkeitszählungen zu bestimmten Zeitpunkten gearbeitet wird (Glasze/Husseini/Mose 2009: 298).

Demnach soll in diesem dritten Schritt entlang der identifizierten und kodierten Diskurselemente sowie ihrer kodierten Verknüpfungen untereinander nach ersten Ordnungs-, Regel- und Klassifikationsprinzipien der narrativen Muster gesucht werden. Die aufgefundenen Charakterisierungen der narrativen Muster können dann auch wieder in Form von Codes dargestellt werden (Diaz-Bone/Schneider 2004: 466f), die nach überlagerten diskursiven Regelprinzipien entlang der Äquivalenz-, Oppositions-, Kausalitäts- und Temporalitätsprozeduren geordnet sind. Auch vorherrschende narrative Muster der Subjekt-, Ereignis und Objektkonstitution können dabei zur Ableitung der diskursiven Regeln dienen.

Die herausgearbeiteten diskursiven Regeln, die das hauptsächliche Resultat der Forschungsarbeit darstellen, werden aus der Analyse der Regelmäßigkeiten abgeleitet. Ebenso ist dieses Vorgehen in erster Linie ein interpretativer Schritt (Glasze/Husseini/Mose 2009: 298), der sich im Wesentlichen auf die Auswertung des Textmaterials und die im Prozess der Kodierung angefertigten Memos zur Erfassung der narrativen Muster stützt. Darüber hinaus können auch im Prozess der Kodierung entstandene integrative Diagramme die Analyse erhellen (Mey/Mruck 2009: 136), da durch sie das Zusammenwirken der Codes und Kernkategorien oder in der Sprache Laclaus und Mouffes die Knotenpunkte und wesentlichen Äquivalenz- sowie Oppositionsbeziehungen zu anderen Elementen, in einem relationalen Raum und je nach Forschung auch entlang einer Zeitachse visualisiert werden können (s. bspw.: Glasze 2013b: 168, 171, 213, 215). Die Ausgangsthese zur Erfassung der diskursiven Regeln ist, dass die Regelmäßigkeiten ein im Text materialisierter Ausdruck der Regeln sind. Jedoch ist es wichtig dabei, wie Glasze et al. betonen, dass trotz allem die Regeln nicht als statisch aufgefasst werden dürfen, sondern »sie sind vielfältig, können widersprüchlich sein und unterliegen fortwährenden Veränderungen.« (Glasze/Husseini/Mose 2009: 298) Auch wenn aus all den Daten ein Ergebnis produziert wird, so ist für die hier an Laclau und Mouffe orientierte Diskursanalyse, die sich auf die Analyse der narrativen Muster eines diskursiven Feldes bezieht, das Gleiche wie für die Grounded Theory zu konstatieren: Die Forschung sowie die Ergebnisse der kodierenden Verfahren müssen als »ein kontinuierlicher und endloser Prozess« begriffen werden, der als niemals abgeschlossen betrachtet werden kann (Mey/Mruck 2009: 136f).

5.3.2.4 Welche Korpora wurden für die Diskursanalyse kodiert?

Um auf Regelmäßigkeiten im Rahmen kodierender Verfahren in Diskursanalysen stoßen zu können, empfiehlt es sich nach Glasze (Glasze 2013b: 116f) mit umfangreichen Textkorpora zu arbeiten. Die Zusammenstellung der Textkorpora zur GB basierte dabei zunächst auf einem vorhandenen Kontextwissen, welches sich durch eine eingängige Literatur-, Datenbank- und Archivrecherche im Feld angeeignet wurde. Hier wurde, wie oben beschrieben, in Anlehnung an das Modell der Global Governance-Architektur nach Dirk Messner (1999: 13) für einige der hier jeweils relevanten Handlungsebenen das Datenmaterial mindestens einer beispielhaften Sprecher*innenposition in den Diskurskorporus eingespeist. Jede Sprecher*innenposition entspricht hierbei einer relevanten und spezifischen Form einer institutionalisierten Gruppe oder Diskurskoalition des Feldes GB. Dieses erste digitale und geschlossene Diskurskorporus bildet dabei einen fallspezifischen und modellhaften Ausschnitt des hegemonial strukturierten Feldes Glo-

baler Bildung in den Global Governance-Strukturen. Bei der Auswahl des Textmaterials wurde dabei darauf geachtet, ein möglichst einheitliches Genre aus einflussreichen Strategie-, Konzept- und Positionspapieren oder bildungstheoretischen und -praktischen Beiträgen zur Gestaltung GB der jeweiligen Sprecherpositionen im zu untersuchenden Zeitraum von 2012 bis 2018 für das diskursive Feld zusammenzustellen.

Da im feinanalytischen Schritt der kodierenden Verfahren noch weitere Detailfragen an das Material gerichtet aufkommen und darüber hinaus noch weitere, hier allen voran marginalisierte Diskursbereiche, in die Analyse des diskursiven Feldes einfließen sollen, ist es im Gegensatz zu den für die lexikometrischen Untersuchungen genutzten geschlossenen Textkorpora sinnvoll, in einer an Michel Pecheaux angelehnten Idee im Rahmen der Analyse narrativer Muster mit einem »offenen [...] sich erweiternden Korpus« zu arbeiten (Glasze 2013b: 117; s. auch Busse 2000: 44), d.h. dieses im Zuge der Analyse je nach Befragung des Materials durch das Kontextwissen zu verkleinern oder zu ergänzen (Glasze 2008: 201). Somit wird auch dem Umstand Rechnung getragen, dass zu Beginn eines Forschungsprozesses nicht alle relevanten Positionen des zu untersuchenden Diskursfeldes überblickt werden können, wie Glasze (2013b: 117) bemerkt. Er schlägt deshalb vor, sich »an die Methode des *theoretical sampling* der *grounded theory* anzulehnen.« (Glasze 2013b: 117) Kern dieser Methode ist es, dass Datenerhebung und -auswertung nicht in unterschiedlichen Phasen geschehen, sondern sukzessiv vollzogen werden. D.h. Erhebung und Auswertung wechseln sich hier stets ab und auf der Basis der ausgewerteten Ergebnisse wird entschieden, welches nächste Material in die Untersuchung mit eingeht. Im Mittelpunkt der Auswahl steht dabei die Relevanz des Materials für die Ergebnisse der Forschung zu einem bestimmten Zeitpunkt (Mey/Mruck 2009: 110). Das Kriterium der Relevanz muss dabei als eines begriffen werden, das das jeweilige Phänomen (hier: die narrativen Muster) am besten zu verstehen und zu erklären hilft. Diese vielen immer fortlaufenden Schritte, werden fortgesetzt, bis beim Prozess des Hinzuziehens und Erhebens weiteren Datenmaterials kein weiterer Erkenntnisgewinn für die Analyse mehr zu erzielen ist (Mey/Mruck 2009: 111f).

5.3.2.5 Die Nutzung von MAXQDA zur Kodierung des Textmaterials

Programme zur computergestützten qualitativen Textanalyse mit Textverwaltungs-, Such- und Kodierfunktionen können die Analyse von narrativen Mustern um ein Vielfaches erleichtern. Seit den 1980er Jahren stehen sogenannte QDAS auch »Qualitative Data Analysis Softwares« zur computergestützten qualitativen Analyse von Texten zur Verfügung. Diese sind dabei nicht als Werkzeuge zur Analyse zu betrachten, sondern vielmehr dienen sie der Strukturierung und Organisation von Textdaten (Diaz-Bone/Schneider 2004: 457). Dabei gibt es einer Vielzahl von Kategorien, in die die QDAS klassifiziert werden können. Die für Diskursanalysen relevanten QDAS werden als »Code-and-Retrieve-Programs« bezeichnet. Diese ermöglichen das Suchen, Finden, Kodieren und Systematisieren von Texten. Mit ihnen besteht die Möglichkeit »Textpassagen schnell und bequem »codieren« und auf solche »codierten« Passagen systematisch zugreifen zu können.« (ebd.: 460) Diese grundlegenden Funktionen erlauben es »eine im Material fundierte und rekursiv verfahrenende Analyse/Interpretation der Daten zu unterstützen.« (ebd.: 460) Außerdem steht für die interpretativen Auswertungen den Nutzer*innen eine Memo-Verwaltung »zum Erstellen und Organisieren von z.B.

Daten-, Code-, Theorie-memos etc.« zur Verfügung (ebd.: 460). Zusätzlich bieten sie die Möglichkeit den gesamten Verlauf und die Ergebnisse der Auswertungen in Form von begrifflichen Netzwerken anschaulich darstellen zu können, um im Nachhinein die aus den Daten entwickelten Codes stets weiterentwickeln und verdichten zu können (ebd.: 460). Zudem können auch die Verkettungen zwischen verschiedenen Codes und Begriffen aufgesucht werden. QDAS bieten damit grundlegend die Möglichkeit, verschiedenste Datenmaterialien für qualitative Forschungen gebündelt zu sammeln und auch die Auswertungsergebnisse transparent nachvollziehbar und darstellbar zu machen (ebd. 460). Diaz-Bone und Schneider betonen jedoch vehement, dass gerade für den diskursanalytischen Einsatz von QDAS die Passungsverhältnisse zwischen den theoretischen Grundprämissen dieser und der Methodik der Diskursanalyse unbedingt angepasst werden müssen. Vor allem, da eine Vielzahl der verfügbaren QDAS »als eine ›in Software geronnene Methodologie‹ der Grounded Theory zu verstehen ist« (ebd.: 463). Für die oben beschriebene Umsetzung des kodierenden Verfahrens wurde sich für die Software MAXQDA entschieden.