

tionen und Deutungsakte der Sprecher*innen gerichtet: Welches Phänomen konstruieren die Akteur*innen in ihren Sprechakten, welche (ursächlichen und intervenierenden) Bedingungen schreiben sie diesem Phänomen zu, was wird als Kontext identifiziert und welche Strategien werden entweder identifiziert oder vorgeschlagen, um auf das Phänomen zu reagieren und welche Konsequenzen werden dem Phänomen zugeordnet oder noch befürchtet? Die Integration der GT in die WDA als übergeordnete methodologische Richtschnur basiert also schlicht auf der Modifizierung der analytischen Perspektive auf das, was Handlung und Feld konstituieren.

3.2 Methodisches Vorgehen

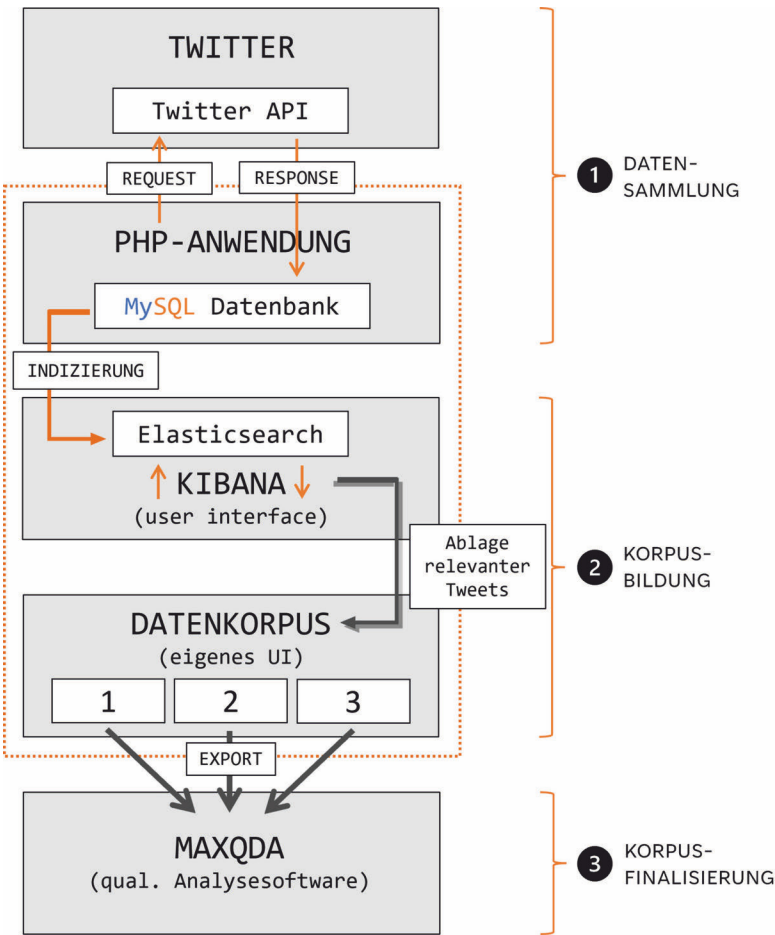
Die Forschungsfrage dieser Arbeit birgt die Schwierigkeit, dass sie zunächst weder auf einen konkreten Zeitraum noch auf ein bestimmtes Ereignis, ein spezifisches Hashtag oder einen ähnlichen begrenzenden Faktor abzielt. Aus diesem Grund wurden hier ereignis- und zeitpunktübergreifende ›Themenfelder‹ in Anlehnung an die drei zentrierten Bündnisgruppen der AfD (*Alternative Homosexuelle*, *Juden in der AfD* und *Neudeutsche Alternative*) definiert, um darüber einen Zugang zu den diskursiven In- und Exklusionsprozessen und der Bearbeitung der damit verbundenen Ambivalenzen zu erhalten.⁹ Um die dafür notwendige Explorativität im methodischen Vorgehen zu ermöglichen und eine vorschnelle Einschränkung der Datenbasis zu vermeiden, wurde unter Einbezug technisierter Verfahren ein mehrstufiges Sampling- und Erhebungsverfahren entwickelt. Der Logik der *Grounded Theory* folgend, ist es eng verschränkt mit der Analyse der Daten. Die genaue Vorgehensweise und die darin eingebetteten Entscheidungen werden nachfolgend schrittweise dargelegt.

3.2.1 Sampling und Datensammlung

Der Samplingprozess lässt sich zur besseren Darstellung in drei Ebenen gliedern, die mit unterschiedlichen Plattformen der Datensammlung, -sortierung und -auswertung korrespondieren: Zunächst wurden (1) Diskursarena (Twitter [X]) und Akteurspektrum (AfD) eingegrenzt und, mithilfe eines *Twitter Developer Accounts* und einer daran angeschlossenen Datenbank, für die kontinuierliche Datensammlung zugänglich gemacht. Auf einer zweiten Ebene (2) wurden mithilfe der auf *Elasticsearch* basierenden Analyseplattform *Kibana* drei an den Interessengruppen ausgerichteten Datensätze erstellt. Eine (3) systematische Analyse erfolgte in der Datenanalysesoftware *MAXQDA*.

9 Für eine genauere Ausführung der Fragestellung und eine vertiefende Erläuterung zur Frage der Widersprüchlichkeit im Kontext der Interessengruppen vgl. Kap. 1.

Abbildung 1: Samplingprozess [eigene Darstellung]



›Zusammengehalten‹ werden die Ebenen durch das *theoretical sampling* als leitendes Prinzip der Fallauswahl. Dadurch ist es – sowohl technisch als auch methodologisch – möglich, den Forschungsprozess mit nur rudimentären Eingrenzungen zu beginnen und sich dem Gegenstand und den relevanten Daten explorativ anzunähern. Als Ausgangspunkt für die sukzessive Konkretisierung des Datenkorpus wurden daher, in Anlehnung an die Fragestellung und eine erste Sondierung des Feldes (vgl. Keller 2011a: 86), lediglich Vorentscheidungen bezüglich des zu untersuchenden Akteur*innenspektrums, der Diskursarena und der sich daraus ergebenden Art des Materials getroffen. Als Diskursarena wurde die Microblogging Plattform *Twitter (X)* gewählt, weshalb das Material Twitter-Posts (Tweets) des Akteur*innenspektrums, bestehend aus AfD-Mitgliedern, Parteiverbänden und -fraktion, der Jugendgruppe der AfD (*Junge Alternative*) sowie parteiinterne Interessengruppen (insb. *Alternative Homosexuelle, Neudeutsche Alternative* und *Juden in der AfD*), umfasst. Im Laufe der Forschung wurden

Datensammlung und -analyse auf Posts beschränkt, die zwischen dem 1. Juli 2015 und dem 31. September 2021 verfasst oder geteilt wurden.¹⁰

3.2.1.1 Diskursarena und Datentyp: Twitter (X)

Mit 14,1 Millionen monatlichen Nutzer*innen im Januar 2023 ist *Twitter (X)* zunächst nur die fünftstärkste Social-Media-Plattform in Deutschland (vgl. DataReportal 2023)¹¹ und auch die AfD kann auf den ersten Blick mit nur 191.893 Follower*innen nicht an ihren Erfolg auf *Facebook* anknüpfen. Auf der in Deutschland meistgenutzten Social-Media-Plattform liegt sie mit 557.678 Follower*innen weit vor den anderen deutschen Parteien.^{12,13} Dennoch erreicht die AfD auch auf *Twitter (X)* die mit Abstand höchste Interaktionsrate (0,62 Prozent)¹⁴ (vgl. Mienert 2021). Die Diskrepanz zwischen der Follower*innenanzahl und der Interaktionsrate lässt sich mit der besonderen Dynamik der Plattform erklären, der ein besonderes »populistisches Kommunikationspotenzial« zugeschrieben wird (vgl. Datts 2020: 75; vgl. Dittrich 2017: 15f.; vgl. Ernst et al. 2017; vgl. Jacobs et al. 2020: 613f.; vgl. Jacobs/Spierings 2019: 1685). Sie ergibt sich aus der, vor allem bis zum Inkrafttreten des *Gesetzes zur Verbesserung der Rechtsdurchsetzung in sozialen Netzwerken* (Netzdurchsetzungsgesetz – NetzDG) im Oktober 2017 nahezu unmoderierten Kommunikation (vgl. Dittrich 2017: 5; vgl. Engesser et al. 2017: 1123). Aber auch die hohe Frequentierung durch nicht-registrierte Nutzer*innen und Journalist*innen, die die Inhalte besonders häufig in traditionelle Medienkanäle übertragen, tragen dazu bei (vgl. Datts 2020: 74; vgl. Klinger/Svensson 2015: 11). *Twitter (X)* wird daher eine beachtliche »significance for the agendas of the mass media and thereby for the wider perception of political issues« (Datts: 2020: 74) zugeschrieben und erfüllt so eine doppelte Funktion für Populist*innen: Zum einen lassen sich die traditionellen Medien in ihrer Funktion als Intermediäre und Gatekeeper umgehen, um direkt »mit dem Volk« zu kommunizieren, zum anderen steigen gleichzeitig die Chancen, die massenmediale Berichterstattung zu beeinflussen – allerdings ohne sich den Regeln der, im Populismus

10 Der Juli 2015 wurde als Startpunkt gewählt, da hier i.d.R. der erste »Rechtsruck« der Partei verortet wird (vgl. Akeel 2021: 303; vgl. Heinze 2020: 41). Entscheidend für die Festlegung des Endpunktes war die einsetzende theoretische Sättigung.

11 Insgesamt nutzen im Januar 2023 rund 70 Millionen Menschen in Deutschland Social Media (vgl. DataReportal 2023).

12 Die Angaben zu den Follower*innenanzahlen auf *Facebook* und *Twitter (X)* stammen aus eigener Recherche, Stand: 23.04.2023.

13 Zwar gibt es Studien, die populistische Kommunikation auf *Twitter (X)* und *Facebook* vergleichend betrachten (z.B. Ernst et al. 2017 und Jacobs et al. 2020), eine aktuelle, systematisch vergleichende Analyse der plattform-spezifischen Kommunikationsstrategien der *Alternative für Deutschland* fehlt bisher jedoch. Da die Plattformen unterschiedliche Topografien (Bossetta 2018) besitzen, die die Kommunikationsdynamik beeinflussen, würde eine solche Untersuchung helfen, in Zukunft Forschungsfragen präziser an die jeweiligen Gegebenheiten der Plattformen und damit des Materials anzupassen bzw. besser das passende Material für unterschiedlich ausgerichtete Forschungsfragen zu wählen.

14 Auf dem zweiten Platz liegt *Die Partei* mit nur noch 0,28 Prozent (Stand: 27.04.2023) (vgl. Mienert 2021). Die Interaktionsrate wird aus der »[durchschnittlichen] Anzahl der Interaktionen (Likes, Kommentare, Shares etc.) pro Tag auf Posts/Beiträge eines Tages im Verhältnis zur Anzahl der Fans/Follower desselben Tages in diesem Jahr« (ebd.) berechnet.

häufig als ›Lügenpresse‹ denunzierten, Massenmedien zu unterwerfen (vgl. Datts 2020: 75; vgl. Dittrich 2017; vgl. Jacobs et al. 2020: 614f.; vgl. Stier et al. 2017: 2; vgl. Van Kessel/Castelein 2016: 595). Trotz der vergleichsweise geringen Nutzer*innenzahl ist es über *Twitter (X)* also möglich, kosten- und ressourcengünstig eine breite Masse an Menschen zu erreichen und zu beeinflussen (vgl. Datts 2020: 76; vgl. Dittrich 2017: 18; vgl. Ernst et al. 2017: 1351f.; vgl. Gil De Zúñiga et al. 2020: 587; vgl. Jacobs/Spierings 2019: 1683f.; vgl. Klinger/Svensson 2015: 6). Auch die Plattformstruktur von *Twitter (X)* begünstigt eine bestimmte Aufmerksamkeitsökonomie. Über *retweets* (RT) oder Quote-Tweets (RT mit eigenem Kommentar, der über dem Original angezeigt wird) und die Verwendung von *hashtags* (#) können die auf 280 Zeichen¹⁵ beschränkten Äußerungen in kürzester Zeit einem breiten Publikum zugänglich gemacht werden. Eine direkte Interaktion ist durch Erwähnungen (*mentions* [@]) anderer Nutzer*innen oder das Antworten auf Tweets (*replies*) möglich.

In dem hier untersuchten Kontext hat sich gezeigt, dass die begrenzte Zeichenzahl zu einer besonders hohen Bedeutungsichte in den Aussagen und zu komplexen Verweisstrukturen führen kann. Die Zeichenbegrenzung zwingt die Autor*innen zu reflektieren, *was* ausgesagt werden soll, aber auch, *wer* die Aussagen verstehen können soll. Das führt unter anderem zur Verwendung bestimmter Ausdrucksweisen, Abkürzungen, politischer Kampfbegriffe oder sogenannter *dog whistles*.¹⁶ Da die hohe Interaktionsrate und große massenmediale Resonanz, die die AfD immer wieder mit ihrer *Twitter*-Präsenz erzeugen kann (vgl. Maurer et al. 2023), darauf verweisen, dass die Partei die plattformeigenen und -übergreifenden Strukturen bewusst und geschickt im Sinne eines ›hybrid media system‹ (Chadwick 2013) nutzt, wurde *Twitter (X)* für diese Untersuchung als Diskursarena ausgewählt. Wie für andere natürliche Daten gilt auch für Tweets: »in these publications, people converse, announce position, argue with a range of eloquence etc.« (Glaser/Strauss 2008: 163).

Aus den Eigenheiten der Plattform ergeben sich allerdings auch einige Schwierigkeiten, die bei der Datenerhebung berücksichtigt werden müssen. Zwar bietet *Twitter (X)* einen vergleichsweise niedrigschwelligen Zugang zu großen Datenmengen (vgl. Mayerl/Faas 2019: 1027ff.), aber die generelle Flüchtigkeit von Online-Inhalten (vgl. Bozdağ 2018: 131), die algorithmenbasierten Filter (vgl. Bossetta 2018: 476f.) und die Einschränkungen durch die Nutzungsbedingungen erschweren einen systematischen Zugriff auf die Daten (vgl. Twitter Inc. 2022a). Sowohl bei einer manuellen Erhebung (z.B. durch Screenshots oder Transkription) als auch beim Einsatz von *browser crawlers* oder anderen *web scraping tools*¹⁷ kann nicht nachvollzogen werden, nach welchen Kriterien zum Beispiel

15 Von 2006 bis 2017 lag die Zeichenbegrenzung bei 140 Zeichen. Der vorliegende Datenkorpus umfasst demnach auch Tweets, die noch dieser ersten Einschränkung unterliegen.

16 Eine ›dog whistle‹ oder ›dog whistling‹ bezeichnet eine Kommunikationsstrategie, bei der gezielt kodierte Sprechakte eingesetzt werden, um diskriminierende, rassistische oder anderweitig gegen soziale Normen verstoßende Inhalte öffentlich an eine ›wissende‹ Zielgruppe zu kommunizieren (vgl. Bhat/Klein 2020: 153f.). Aufgrund ihres »potential for deniability« (ebd.: 154) hat sich diese Form der Kommunikation in den letzten Jahren verstärkt in der *Alt-Right*, der *Neuen Rechten* und rechtspopulistischen Gruppen etabliert – so auch bei der AfD (vgl. Pfahl-Traughber 2020: 90).

17 *Web scraping* ist eine automatisierte Form der Datensammlung, mit der es möglich ist, große Datenmengen aus Webseiten zu extrahieren. In vielen Kontexten fällt *web scraping* in eine rechtliche

Suchergebnisse durch den Twitter-Algorithmus vorgefiltert wurden. Ebenso wenig kann ausgeschlossen werden, dass identische Suchparameter bei verschiedenen Suchdurchgängen jedes Mal unterschiedliche Ergebnisse erzeugen. Das für diese Arbeit entwickelte Samplingverfahren hat daher zum Ziel, die benannten Schwierigkeiten handhabbar und die Daten, in Form eines konsistenten und durchsuchbaren Datenkorpus, für eine qualitative Analyse zugänglich zu machen.

3.2.1.2 Datensammlung

Eine Möglichkeit, direkten und »breiten Zugriff auf öffentliche Twitter Daten [zu erhalten], die Nutzer wissentlich mit der Welt geteilt haben« (Twitter Inc. o.J.), bieten die Datenschnittstellen von Twitter (X), sogenannte Applikationsprogrammierschnittstellen (API). Auf sie kann über *Twitter Developer Accounts* zugegriffen werden: eine Erweiterung, die vor allem für Unternehmen, Wissenschaftler*innen und Entwickler*innen gedacht ist, aber prinzipiell für jeden vorhandenen Account freigeschaltet werden kann.¹⁸ Darin können zum Beispiel unterschiedliche Projekte zur Datenanalyse angelegt oder externe Anwendungen zur Aufbereitung der Daten an die API angeschlossen werden (z.B. eine integrierte Entwicklungsumgebung (IDE) wie *R Studio* oder eine Analysesoftware wie *MAXQDA*). Für die vorliegende Untersuchung wurde eine (relationale) MySQL-Datenbank angelegt, in der mit einer PHP-Anwendung die über die API abgerufenen Daten abgelegt wurden.¹⁹ Diese Vorgehensweise sollte einen kontrollierten und uneingeschränkten Zugriff auf eine konsistente Datenbasis ermöglichen, ohne im Vorhinein zu viele Eingrenzungen vornehmen zu müssen und möglichst offen vorgehen zu können. Ziel war es, eine Art Bibliothek anzulegen, die einerseits erlaubt, die Fallauswahl nach der Logik des *theoretical sampling* aufrechtzuerhalten und andererseits die forschungsökonomische und transparente Sicherung der Daten zu gewährleisten.

Grauzone. Es ist nicht per se illegal, aber unter bestimmten Umständen zivilrechtlich verfolgbar (vgl. Phoenix 2022).

- 18 Einen Twitteraccount als *Twitter Developer Account* freischalten zu lassen, stand zum Zeitpunkt der Erhebung prinzipiell jeder Person offen. Die Freischaltung erfolgte nach einer Auskunft über den Nutzungszweck und der Beantwortung eventueller Nachfragen (vgl. Twitter Inc. 2022b). Im Zuge der Übernahme der Plattform durch Elon Musk wurden die Zugänge im Frühjahr 2023 allerdings vorübergehend geschlossen und nach Auskunft von *Twitter (X)* überarbeitet. Die anschließend eingeführten kostenpflichtigen Zugänge ermöglichten dann nur noch z.T. stark eingeschränkten Datenzugriff zu sehr hohen monatlichen Kosten (z.B. 100 USD im Monat für den Zugriff auf monatlich maximal 10.000 Tweets im Vergleich zu den vorherigen 10 Millionen). Umfragen der *Coalition for Independent Technology Research* (2023) und *Reuters* (Dang 2023) weisen darauf hin, dass im Zuge dessen Ende 2023 mehrere hundert Forschungsprojekte entweder Gefahr liefen, komplett zum Erliegen zu kommen, oder bereits gestoppt werden mussten. Ebenfalls betroffen waren freie Tools, wie das *Botometer*, Erweiterungen zur Barrierefreiheit oder das sog. *Disaster-Mapping*. Seit dem Inkrafttreten des *Digital Service Act* im Februar 2024 für alle sog. VLOPs (= Very Large Online Platform) hat *Twitter (X)* erneut einen Forschungszugang ermöglicht. Allerdings bestehen weiterhin Zweifel an der Transparenz bezüglich der zugänglichen Daten. Neben der veränderten »Plattformkultur«, die Auswirkungen auf die Eignung der Daten und mögliche wissenschaftlichen Fragestellungen hat, ist vielen Forscher*innen darüber hinaus das Risiko einer rechtlichen Verfolgung durch die Plattform zu groß (vgl. Reuters 2023; vgl. auch Ledford 2023, Murtfeldt et al. 2024).
- 19 Dafür wurde die Anwendung CITCAT verwendet (Adams 2023).

Als Grundlage der Bibliothek bzw. Ausgangspunkt für das Sample wurden Accounts aus dem oben benannten Akteur*innenspektrum festgelegt und ihnen »gefolgt«. Auf diese Weise konnten Tweets dieser Accounts sowohl retrospektiv als auch realzeitlich gesammelt werden. Dafür wurden mit der PHP-Anwendung minütlich Suchanfragen (*request*) an verschiedene Endpunkte (»friends/list«, »search/tweet«, »statuses/home_timeline«) der Twitter-API gestellt und die Ergebnisse (*response*) in der Datenbank abgelegt (Abbildung 1).

Auswahl der Accounts

Die initiale Auswahl der Accounts kann anhand der drei Modi des offenen Samplings – gezielt, systematisch und zufällig – beschrieben werden (vgl. Strauss/Corbin 1996: 155f.). Zunächst wurden gezielt Accounts von bekannten AfD Politiker*innen, Fraktionen und Gruppierungen ausgewählt (Tabelle 1).

Tabelle 1: Accounts nach Gliederung und Organisationseinheit²⁰

Organisa- tionseinheit	EU	Bund	Land	Kreis	Ort/ Bezirk	sonst.	ges.
Einzelpersonen	3	36	30	5	6	1	81
Parteiverbände & -fraktionen	1	3	22	43	4	0	73
Junge Alternative	0	1	3	0	4	0	8
Interessenverbände & sonstige	0	6	1	0	0	0	7
<i>gesamt</i>	4	46	56	48	14	1	169

Dabei wurde darauf geachtet, Accounts aus allen Gliederungsebenen (Stadt, Kreis, Land, Bund) zu wählen und möglichst viele Bundesländer und Regionen abzudecken. Der weitere Auswahlprozess oszillierte zwischen Systematik und Zufälligkeit. Einerseits wurden die ausgewählten Accounts systematisch nach Verweisen auf andere relevante Accounts durchsucht, andererseits wurden Accounts gewählt, die der Twitter-Algorithmus vorschlug. Die genaue Zusammensetzung der Accounts hat sich im Verlauf des Forschungsprozesses weiter verändert, indem neue Accounts, die sich im Laufe der Analyse als relevant erwiesen, hinzukamen, während andere entfernt wur-

20 Eine vollständige Auflistung aller Accounts findet sich im Anhang 1.1.

den.²¹²² Schlussendlich umfasst der Datensatz 169 Accounts von denen für den Zeitraum vom 1. Juli 2015 bis zum 30. September 2021 insgesamt 533.151 Tweets gesichert wurden.

3.2.1.3 Korpusbildung

Aus diesem umfassenden Datensatz konnten in einem nächsten Schritt gezielt(er) Daten ausgewählt werden, die Auskunft zur Forschungsfrage geben. Für die Exploration solcher natürlichen Daten schlagen Glaser und Strauss vor, »to cultivate functional synonyms for [the researchers] topic in order to explore relevant categories fully« (Glaser/Strauss 2008: 171).²³ Solche »funktionalen Synonyme« werden auf Basis von theoretischer Sensibilität, also theoretischem, feld- und alltagsspezifischem Vorwissen, entwickelt und im Prozess der Forschung kontinuierlich geschärft oder direkt den Daten entnommen (vgl. Truschkat et al. 2005). Für die vorliegende Arbeit bedeutete das, Stichworte zu finden, die solche Daten generieren, die Auskunft über die Konstruktion von Zugehörigkeit und Fremdheit in unterschiedlichen Kontexten und die Aushandlung damit verbundener Widersprüche geben. Dazu wurden die Stichworte thematisch an den drei Gruppierungen (*Alternative Homosexuelle*, *Neudeutsche Alternative*, *Juden in der AfD*) ausgerichtet bzw. davon abgeleitet und drei separate Datensätze erstellt, um mögliche Unterschiede und Eigenheiten der Gruppen und damit assoziierten Gesellschaftsbereichen herausarbeiten zu können. Darüber wurde sowohl der Umgang mit den Interessenvereinigungen seitens der Partei erfasst als auch ostentative und verdeckte Einstellungen gegenüber den von ihnen vertretenen gesellschaftlichen Gruppen (Homosexuelle, Jüdinnen*Juden, Migrant*innen und nicht-Weiße Menschen) offengelegt. Daher umfassen die Stichworte neutrale Deskriptoren für diese Gruppen, aber auch Pejorative, Eigenbezeichnungen oder solche, die auf öffentliche Debatten verweisen, die die respektiven Gruppen betreffen. Einige Beispiele wurden in Tabelle 2 zusammengetragen.²⁴

-
- 21 Das betrifft vor allem Accounts, deren Inhaber*innen keine AfD-Mitglieder sind. Zu Beginn der Analyse waren einige parteilose rechtspopulistische Akteur*innen sowie Medienkanäle (z. B. Birgit Kelle, Ken FM oder COMPACT Magazin) in den Datensatz aufgenommen, im Laufe der Analyse aber wieder verworfen worden. Die einzige Ausnahme stellt Erika Steinbach dar, da sie bereits vor ihrem Parteieintritt 2022 offen für die AfD geworben hat. Accounts von Personen, die im Untersuchungszeitraum in die AfD ein- oder ausgetreten sind, wurden nur für die Dauer ihrer Partei-Mitgliedschaft untersucht.
 - 22 Alle Accounts wurden unabhängig von dem Zeitpunkt, an dem sie in das Sample aufgenommen wurden, nach den gleichen Kriterien durchsucht.
 - 23 Glaser und Strauss beziehen sich in ihrer Diskussion natürlicher Daten vor allem auf den Umgang mit Bibliotheksmaterial (vgl. 2008: 170f., 176ff.). Die Ausführungen dazu können aber in vielfacher Weise auf die Arbeit mit einer Datenbank, als einer Form der Bibliothek, übertragen werden.
 - 24 Eine vollständige Auflistung aller Stichworte sowie die Anzahl der jeweils in der Datenbank generierten Treffer findet sich in Anhang 1.2.

Tabelle 2: Auszug aus den verwendeten Stichworten²⁵

Datensatz 1	Datensatz 2	Datensatz 3
Themenfeld ›Homosexualität‹	Themenfeld ›Jüdisch-sein‹	Themenfeld ›Fremdheit‹
Regenbogen	antisemit*	Ausländer
-flagge	Israel	Bevölkerungsaustausch
-fahne	Judenhass*	Diversity
"Ehe für Alle"	Holocaust	Europa
homo*	Jüd*	Heimat
Frühsexualisierung	Knobloch	Migrant*
gleichgeschlechtlich*	Schuld kult	N-Wort
Kernfamilie	Zentralrat	Rass*

Um die Datenbank möglichst effizient mit den Stichworten durchsuchen zu können, wurden die Daten in *Elasticsearch* indiziert, d.h. die in den Daten enthaltenen Informationen wurden nach bestimmten Kriterien in einem Suchindex katalogisiert. Über die browserbasierte Analyseplattform *Kibana* konnte dann gezielt auf diese Informationen zugegriffen werden, zum Beispiel durch lexikalische Suche (Stichworte), die Eingrenzung nach Datum, *username* oder ähnlichen Faktoren. Die durch die Stichworte erzeugten Treffer wurden soweit möglich²⁶ vollständig gelesen und die für die Fragestellung relevanten Tweets erneut in der Datenbank, diesmal in drei separaten Datensätzen abgelegt. Dieser Vorgang warf auf unterschiedliche Weisen neue Stichworte ab, mit denen der zuvor erläuterte Prozess wiederholt wurde: entweder dadurch, dass in den Ergebnissen direkt interessante Begriffe auftauchten oder aus den ersten parallelen Analyseprozessen Suchhypothesen entstanden, für die neue ›funktionale Synonyme‹ gebildet wur-

25 Die Sonderzeichen (sog. Suchmodifizierer) zeigen bestimmte Manipulationen der Suchfunktion an. Die Schreibweise in Anführungszeichen (") bewirkt eine zusammenhängende Suche aller eingeschlossenen Begriffe. Der Asterisk (*) vor oder nach einem Wortfragment verweist auf eine sog. ›wildcard search‹, bei der dann alle Iterationen dieses Fragments gefunden werden können. Das Suchwort ›Uni*‹ würde in diesem Fall alle Ergebnisse zu ›Uni‹, ›Universität‹, ›Universitäten‹, ›universitär‹, ›universell‹, ›universal‹ usw. generieren und akkumulieren. Durch das Hinzufügen einer Tilde (), der sog. ›fuzzy search‹, lassen sich Wörter mit geringen Abweichungen, die durch Tipp- oder Schreibfehler sowie uneinheitliche Schreibweisen entstehen, mit einbeziehen.

26 Bei einigen Begriffen lag eine so große Menge an Ergebnissen vor, dass gegen Ende der Forschungsperiode nicht mehr alle gelesen werden konnten und dies auch nicht mehr notwendig war. Zur besseren Handhabung wurden daher Zufallsstichproben in einem Umfang von 50–200 Tweets gezogen (die Stichworte, die dies betrifft, sind im Anhang markiert). Glaser und Strauss argumentieren in *The Discovery of Grounded Theory* zwar, dass das statistische Sampling nicht mit der *Grounded Theory* vereinbar sei (vgl. Glaser/Strauss 2008: 63), allerdings handelt es sich im vorliegenden Fall nicht um ein repräsentatives Sampling zur Bildung des Datenkorpus. Vielmehr wurde die Stichprobe aus forschungspragmatischen Gründen gewählt, um Einblick in Teile der Daten zu erhalten und zu bewerten, ob sich eine weitere Auseinandersetzung damit lohnt. Daher lässt sich dieses Vorgehen aus meiner Sicht mit den Prinzipien des *theoretical samplings* vereinbaren und folgt dem Diktum der gegenstandsangemessenen Flexibilität, die Glaser und Strauss über ›formale‹ Prinzipien der Forschung stellen.

den. Vorübergehend abgeschlossen war dieser Samplingschritt, als neue Begriffe keine neuen Inhalte mehr zu eröffnen schienen oder sich einzelne Tweets vermehrt wiederholten. In einem letzten Schritt wurden die drei Datensätze separat aus der Datenbank exportiert und für eine systematische und fokussierte Kodierung in die Analysesoftware MAXQDA importiert.

Für den Samplingprozess ist diese Ebene weiter relevant, weil das *theoretical sampling* auch die Suchbewegungen innerhalb eines abgeschlossenen Datensatzes zur Sättigung und Kontrastierung der Kategorien miteinschließt (vgl. Strauss/Corbin 1996: 152; vgl. Truschkat et al. 2005), aber auch, weil eine durch das Kodieren informierte Rückkehr zu den ersten beiden Ebenen des Sampling stattgefunden hat. Schlussendlich wurden aus den 533.151 Tweets 146.674 für thematisch relevant erachtet und gelesen. Davon wiederum wurden 7314 feinanalytisch ausgewertet und weiter reduziert,²⁷ bis 5602 kodierte Tweets verblieben.

Tabelle 3: Gesamtdarstellung des Samples

Aspekte des Feldes	Diskurs-arena	Akteur*innenspektrum		Untersuchungszeitraum	
Eingrenzung	Twitter (X)	Alternative für Deutschland		01.07.2015 – 31.09.2021	
Korpus	169 Accounts				
	533.151 Tweets				
		Datensatz 1	Datensatz 2	Datensatz 3	gesamt
	Kibana	36.563	15.584	94.527	146.674
	Datensätze	1095	1769	4450	7314
	MAXQDA	772	1231	3599	5602

3.2.2
Datenanalyse

Die Auswertung dieser Daten basiert auf den drei Kodierschritten der *Grounded Theory* (offenes, axiales und selektives Kodieren [Strauss/Corbin 1996]), die durch die Logik des *theoretical sampling* und der theoretischen Sättigung bereits eng mit der Datensammlung verknüpft sind und aus der Forschungsperspektive der *Wissenssoziologischen Diskursanalyse* die Feinanalyse bilden. Die Rekonstruktion der diskurs- und wissensanalytischen Heuristiken, die das Interpretationsrepertoire des Diskurses (Deutungsmuster, Phänomenstruktur, narrative Struktur, Subjektpositionen) – »[d]as typisierende Ensemble von Deutungsbausteinen, aus denen ein Diskurs besteht und das in einzelnen Äußerungen mehr oder weniger umfassend aktualisiert wird« (Keller 2011c: 235) – bilden, geht also aus dieser Analyse hervor und ergänzt sie um die entscheidende Diskursperspektive und ordnende Struktur (*story line*). Der Kodierprozess wird folgend, trotz der starken Verschränkung der Kodierarten in der Praxis, erneut linear dargestellt. Die Beschreibung

27 Redundante Daten wurden im Verlauf der Kodierung aus MAXQDA entfernt.

der Kodierschritte und die damit verbundene Rekonstruktion der Elemente des Interpretationsrepertoires bezieht sich auf die Auswertung eines einzelnen Datensatzes; das Verfahren wurde demnach dreimal durchgeführt. Erst für das Fazit wurden die Ergebnisse der drei Datensätze endgültig zusammengeführt, um erneut übergeordnete Deutungsmuster, gemeinsame Subjektpositionen, Überschneidungen der Phänomenstrukturen und das verbindende Narrativ herauszuarbeiten, mit dem schlussendlich die Forschungsfrage beantwortet werden konnte.

Die Analyse begann mit einem ersten Lesen der Ergebnisse der Stichwortsuche (s.o.), also damit, einen Überblick *über* bzw. Einblick *in* das Material zu gewinnen, um den Erhebungsprozess weiter anleiten zu können. In ersten Memos wurden provisorische Kodes entworfen, um die systematische Kodierarbeit in MAXQDA vorzubereiten und eine vorläufige Sättigung des Materials festzustellen. Diese vorläufige Sättigung war Bedingung für die Überführung der Daten in MAXQDA und die damit einhergehende vorläufige Schließung des Datensatzes.

Der erste Durchgang war vor allem auf das *offene Kodieren* und somit die Aufgabe gerichtet, das Material ›aufzubrechen‹ und möglichst breit zu erfassen, welche Phänomene und diskursiven Konstruktionen von Phänomenen im jeweiligen Teil-Korpus verhandelt werden. Konkret bedeutete das, Diskursfragmente, also einzelne Tweets, Zeile für Zeile zu lesen und die Äußerungen auf Aussagegehalte, referenzierte Ereignisse und Personen etc. hin zu untersuchen, die dann mithilfe von präliminären Kodes oder Phrasen (z.B. ›Darstellung Islam als homophob‹, ›Grundgesetz‹, ›Familienbild‹ ...) festgehalten wurden. Die Benennung oder Umschreibung verlief in diesem Schritt völlig offen, folgte also noch keinem bestimmten System oder Fokus, und erstreckte sich auf unterschiedliche Ebenen der Abstraktion. Im weiteren Verlauf des Analyseschritts erfolgte eine sukzessive Gruppierung und Zuordnung von Kodes zueinander und somit eine Annäherung an die Rekonstruktion diskursiver Elemente. Auch hier ging es aber zunächst vor allem darum, Ordnung herzustellen und Inhalte zu systematisieren. Zusätzlich wurden hier mithilfe von Farbmarkierungen rhetorische und andere Auffälligkeiten festgehalten und erste Textstellen markiert, die potenziell ausschlaggebend für die Rekonstruktion wissensanalytischer Elemente des Interpretationsrepertoires sein würden. Mithilfe der Perspektive des *axialen Kodierens* wurde diese Vielfalt von Inhalten und daran anzuknüpfende Analyseperspektiven, die das ›Aufbrechen‹ des Materials erzeugt hat, auf systematische Weise geordnet. Das geschah einerseits über die Erfassung und Benennung der im Material gefundenen Phänomene als Kategorien und Subkategorien, die mithilfe von Hypothesen in Beziehung zueinander gesetzt und anhand des modifizierten paradigmatischen Modells (vgl. Kap. 3.1) (vgl. Strauss 1998: 55–65; vgl. Strauss/Corbin 1996: 75–86) verdichtet wurden. Dabei wechseln sich Prozesse von Induktion, Abduktion und Deduktion ab und tragen zur Verdichtung, Verifikation, Verwerfung und Neuentdeckung von Erkenntnissen bei. Einige Kodes wurden an dieser Stelle nicht in Kategorien im Sinne des axialen Kodierens überführt, sondern nach stärker subsumptionslogischen Kriterien geordnet und im Rahmen der Rekonstruktion der narrativen Struktur kontextualisierend eingesetzt. Diese wird im letzten Schritt des *selektiven Kodierens*, bei dem es um die finale Zusammenführung der vorherigen Auswertungsbemühungen sowie die finale Rekonstruktion des zentralen Phänomens und dessen Strukturelementen geht, erarbeitet. Hierzu wurde in Anlehnung an Strauss und Corbin (1996: 101f.) erneut das

modifizierte Kodierparadigma verwendet, um die Verbindungen der axialen Kategorien und Deutungsmuster herauszuarbeiten. Ergebnis der Verknüpfungen ist einerseits die *story line* oder narrative Struktur, die die Basis für die Ergebnisdarlegungen in Kapitel 4 bildet. Andererseits lassen sich so die Elemente der zentralen Phänomenstruktur, wie Keller sie vorschlägt, herausarbeiten. Ihre Dimensionen überschneiden sich zum Teil mit Elementen des modifizierten Kodierparadigmas (z.B. Handlungsbedarf und Strategien) und ergänzen sie mit der Differenzierung von Ursachen und Verantwortungszuschreibung. Ziel des selektiven Kodierens ist die begründete, kohärente Integration aller Elemente zu einer *grounded theory*.

Die fertigen ›Kodesysteme‹, die für jeden Datensatz erstellt wurden, sind somit immer eine Fusion der unterschiedlichen Analyseschritte aus *Grounded Theory* und *Wissenssoziologischer Diskursanalyse* und bilden die Grundlage für die hieran anschließende Darstellung der Ergebnisse (Kap. 4) sowie die schlussendliche Interpretation in Kapitel 5. Für diese letzte Zusammenführung wurde kein strukturierter Kodierprozess mehr angewendet. Die Zwischenergebnisse wurden dazu *ähnlich* den Kodierv Verfahren verglichen, wobei der Fokus auf dem Abgleich der verschiedenen Phänomenstrukturen und Schlüsselkategorien miteinander lag. Ziel war es, einerseits die zentralen, und alle Datensätze verbindenden Narrative und übergreifende Schlüsselkategorien herauszuarbeiten sowie andererseits Datenmaterial und Fragestellung noch einmal final auf ihre Passung und Vollständigkeit hin zu überprüfen.