

Den Diskursraum vermessen

*Der Bubble-Stärke-Index (BSI): Ein Instrument zur Diagnose kommunikativer
Geschlossenheit von Diskursräumen*

Forschungsbericht

Seminar „Bursting the Bubble“

Prof. Dr. Hohnsträter & Prof. Dr. Blümelhuber

Universität der Künste Berlin · Gesellschafts- und Wirtschaftskommunikation (GWK)

Sommersemester 2026

Vorgelegt von: Paul Étienne Geißler, Charlotte Krupp und Ben Zarniko

Inhalt

1. Einleitung	3
2. Arbeitsdefinition Bubble	3
3. Messmethode: Das Vier-Dimensionen-Modell	4
4. Das Tool: Der Bubble-Stärke-Index	5
4.1 Das methodische Grundprinzip	5
4.2 Operationalisierung	6
4.3 Funktionsweise	8
5. Was messbar ist, wird (an)greifbar	9
6. Reflexion und Fazit	9
Literatur	11

1. Einleitung

„Das ist halt seine Bubble.“ Kaum eine Diagnose fällt im Alltag so leicht. Der Begriff ist überall, in Talkshows, in Kommentarspalten, im Freundeskreis. Aber sobald man genauer hinschaut, wird er unscharf. Meint Bubble die algorithmische Vorsortierung eines Feeds? Das Leben unter Gleichgesinnten? Oder nur das Gefühl, dass die anderen in einer eigenen Realität leben? Ein Wort, das so vieles heißen kann, bestimmt am Ende nichts. Wer mit Bubbles wissenschaftlich arbeiten will, muss den Begriff deshalb zuerst klären.

An genau dieser Unschärfe ist unsere ursprüngliche Projektidee gescheitert, und aus diesem Scheitern ist die vorliegende Arbeit entstanden. Geplant war ein aktivistisches Werkzeug: ein Tool, das in digitalen Foren transparent deklarierte Gegenrede erzeugt. Es sollte Bubbles automatisiert erkennen und dann aufbrechen. Beim Konzipieren merkten wir, dass schon der erste Schritt nicht funktioniert. Es gibt keine mathematische oder logische Funktion und keine Berechnungsgrundlage, mit der ein Programm entscheiden könnte, ob ein Diskursraum eine Blase ist. Und ohne dieses Erkennen bleibt jede Intervention blind. Sie kann weder begründen, wo sie eingreift, noch feststellen, ob sie etwas verändert hat.

Also haben wir das vorgelagerte Problem zum eigentlichen Projekt gemacht. Aus dem Interventionswerkzeug wurde ein Messinstrument: der *Bubble-Stärke-Index (BSI)*. Kapitel 2 klärt zunächst den Begriff der Bubble. Kapitel 3 begründet unser Messmodell, Kapitel 4 zeigt das Tool, Kapitel 5 diskutiert, was Messbarkeit möglich macht, und Kapitel 6 reflektiert die Grenzen.

2. Arbeitsdefinition Bubble

Bevor wir eine eigene Definition vorschlagen, lohnt ein kurzer Blick auf die bestehenden Begriffe. Die Forschung kennt vor allem zwei. Die *Filterblase* nach Eli Pariser (2011) meint die algorithmisch personalisierte Informationsumgebung des Einzelnen. Die *Echokammer* nach Cass Sunstein (2017) meint die selbstgewählte Umgebung Gleichgesinnter, in der sich Positionen wechselseitig bestätigen. Beide Begriffe beschreiben, wie geschlossene Räume entstehen. Genau deshalb helfen sie uns nur begrenzt. Unser Werkzeug muss an einem konkreten Raum erkennen können, ob er geschlossen ist, egal wodurch die Schließung entstanden ist. Wir brauchen also eine Definition, die nicht bei der Ursache ansetzt, sondern beim beobachtbaren Verhalten.

Ein Wort müssen wir dabei vorab klären: der Diskursraum. Wer von Diskursen spricht, landet schnell bei Michel Foucault (1991). Bei ihm ist ein Diskurs kein konkreter Ort, sondern eine Ordnung, die regelt, was überhaupt sagbar ist. Und diese Ordnung beobachtet niemand von außen, denn jeder, der spricht, steht selbst in ihr. Für ein Messinstrument funktioniert dieser Begriff pragmatisch nicht, weil messen nur kann, wer einen abgegrenzten Gegenstand vor sich hat. Wir verwenden Diskursraum deshalb enger und technischer: ein konkreter Ort im Netz, an

dem viele Menschen öffentlich aufeinander reagieren, etwa eine Kommentarspalte oder ein Forum.

Eine Bubble ist für uns dann die *kommunikative Geschlossenheit eines solchen Diskursraums*: das Ausmaß, in dem der Raum Widerspruch nicht verarbeitet, sondern wegfiltert. In der Sprache der Systemtheorie ist das der Extremfall operativer Schließung: Der Raum behandelt Widerspruch nicht mehr als Information, sondern als Störgeräusch (Luhmann 1984). Eine Bubble ist also kein Ort und keine Gruppe, sondern ein Verarbeitungsstil. Und dieser Stil ist beobachtbar, denn er hinterlässt Spuren im Text: darin, wie ein Raum mit Gegenpositionen, mit Mehrdeutigkeit und mit fremden Quellen umgeht. Genau diese Spuren macht sich unser Messmodell zunutze.

Drei Klarstellungen gehören zu dieser Definition. Eine Bubble ist nicht zwingend politisch; auch eine Fangemeinde oder ein Fachforum kann geschlossen sein. Einigkeit allein macht noch keine Blase; entscheidend ist der Umgang mit abweichenden Stimmen. Und Geschlossenheit ist graduell; deshalb bauen wir einen Index und vergeben kein Etikett.

3. Messmethode: Das Vier-Dimensionen-Modell

3.1 Ein Ansatz unter vielen

Eine Definition macht noch keine Messung. Kommunikative Geschlossenheit ist ein Konstrukt, und man kann sie auf viele Arten messen: über Netzwerke, über die Vielfalt von Feeds, über Meinungslandschaften. Wir setzen am sprachlichen Verhalten an, weil sich ein Verarbeitungsstil im Kommunizieren selbst zeigt und weil Kommunikation öffentlich beobachtbar ist, ohne Zugriff auf das Innenleben der Plattformen.

Die vier Dimensionen folgen dem Weg der Kommunikation durch einen Raum: Ein Inhalt muss ankommen (D1), jemand spricht (D2), der Raum verarbeitet das Gesagte (D3), und er spricht in einer eigenen Sprache nach innen und außen (D4). An jeder dieser Stationen kann sich ein Raum schließen, und hinter jeder steht eine eigene Forschungsrichtung, von der algorithmischen Personalisierung über Echokammern und Gruppenpolarisierung bis zur deliberativen Öffentlichkeit nach Jürgen Habermas (1981), in der das bessere Argument zählt und nicht die lautere Stimme.

Einen Anspruch auf Vollständigkeit erheben wir nicht. Ein etablierter Standard, welcher Indikator eine Bubble am zuverlässigsten anzeigt, fehlt bislang. Unsere Auswahl ist deshalb ein begründeter Vorschlag, und Dimensionen lassen sich austauschen oder ergänzen, denn eine neue Dimension braucht nur eine Leitfrage, ein Label und eine Formel (Kapitel 4.2).

3.2 D1: Algorithmische Selektion

D1 erfasst, wie stark vorsortiert ist, was ein Raum überhaupt zu sehen bekommt. Das ist die Filterblase im engeren Sinn. Empfehlungssysteme optimieren auf Verweildauer und nicht auf

Erkenntnis; die Verengung der Zufuhr ist also kein Unfall, sondern Geschäftsmodell. D1 gehört ins Modell, weil es die technische Vorstufe aller anderen Dimensionen ist. Was nie ankommt, kann weder verarbeitet noch beantwortet werden.

3.3 D2: Soziale Homophilie

D2 misst die positionale Einseitigkeit eines Raums: Wie stark ziehen die Stimmen in dieselbe Richtung, gemessen an einer Bezugsposition, also an der Aussage, auf die alle reagieren? Menschen umgeben sich bevorzugt mit Ähnlichen, und in Gruppen Gleichgesinnter verstärken sich Positionen weiter. Dieses Muster heißt Homophilie und gibt der Dimension ihren Namen. D2 gehört ins Modell, weil die Zusammensetzung eines Raums bestimmt, ob Widerspruch überhaupt eine Andockstelle findet. Hinreichend ist sie nicht, denn Einigkeit ist noch keine Blase (Kapitel 2).

3.4 D3: Kognitive Geschlossenheit

D3 ist der Kern des Instruments, denn hier wird die Arbeitsdefinition unmittelbar gemessen: der Umgang des Raums mit Widerspruch und Mehrdeutigkeit. Vier Teilaspekte machen das konkret. Werden fremde Quellen anders abgewertet als eigene? Wie dicht sind pauschal abwertende Begriffe, sogenannte Kampfbegriffe? Wird Mehrdeutigkeit ausgehalten oder zwanghaft aufgelöst? Und werden Gegenargumente inhaltlich bearbeitet oder reflexhaft abgewiesen? In dieser Dimension wird die Unterscheidung zwischen Wir und Die zur Verarbeitungsregel: Sie entscheidet, was in einem Raum überhaupt noch als Argument gilt.

3.5 D4: Semantische Verständlichkeit

D4 prüft die Sprache eines Raums als seine Außengrenze. Wie hoch ist der Anteil an Gruppenjargon? Wie viele Äußerungen sind ohne Insider-Vorwissen unverständlich? D4 gehört ins Modell, weil ein Raum intern respektvoll diskutieren und trotzdem nach außen abgeschottet sein kann. Die Dimension misst diese Anschlussfähigkeit nach außen, unabhängig davon, wie innen gesprochen wird.

4. Das Tool: Der Bubble-Stärke-Index

Das entwickelte Werkzeug ist eine lauffähige Webanwendung, die sich dieses Modells bedient. Sie verarbeitet einen Diskursraum von der Rohquelle bis zum interpretierten Befund. Zwei Grundsatzentscheidungen prägen das Instrument.

4.1 Das methodische Grundprinzip

Warum es keinen objektiven Score gibt

Der BSI liefert bewusst keinen „objektiven“ Einzelwert. Geschlossenheit ist ein Konstrukt ohne Nullpunkt und ohne absolute Obergrenze; jede Zahl ist nur innerhalb einer offengelegten Messvorschrift sinnvoll. Außerdem messen die vier Dimensionen Verschiedenes: Ein Raum kann sozial sehr homogen und zugleich kognitiv völlig offen sein (Kapitel 4.2), und eine

einzelne Gesamtzahl würde das unsichtbar machen. Der BSI arbeitet deshalb mit Anteilswerten zwischen 0 und 1 je Dimension, die als Profil gelesen und relativ zu einer gemessenen Referenz eingeordnet werden.

Label-dann-Score

Das methodische Rückgrat ist das Prinzip *Label-dann-Score*. Die KI wird ausschließlich zum Etikettieren eingesetzt: Sie entscheidet zum Beispiel, ob eine Äußerung einen Kampfbegriff enthält. Die Kennzahlen berechnet sie nie selbst; die entstehen aus den Labels nach den offengelegten Formeln in Tabelle 1. Das Modell arbeitet deterministisch und kann keine Zahl erfinden. Zusätzlich bleibt der Mensch in der Schleife: Jedes Label ist in einer Prüfansicht einsehbar und per Klick korrigierbar, woraufhin der Wert neu gerechnet wird. Eine systematische Prüfung, wie oft Mensch und Modell übereinstimmen, steht noch aus; bisher zeigen stichprobenartige Kontrollen in dieser Prüfansicht plausible Labels.

Geltungsbereich: wo die Methodik greift

Im Kern ist die Methodik ein Kodierschema und damit eine automatisierte Form der strukturierenden Inhaltsanalyse (Mayring 2015): Man kann jede Leitfrage aus Tabelle 1 auch von Hand beantworten und Anteile ausrechnen; das Tool automatisiert das nur. Daraus folgt der Geltungsbereich. Die Messung braucht Räume mit vielen eigenständigen Stimmen, die aufeinander oder auf einen gemeinsamen Reiz reagieren, also Kommentarspalten, Foren und Threads. An einem Einzeltext wie einem Roman greift sie nicht, denn gemessen wird das Verhältnis vieler Äußerungen zueinander. Wir messen exemplarisch YouTube-Kommentarspalten, und zwar relativ zum jeweiligen Video: Ein Raum, der sich zustimmend um die Videoposition schart, ist ein Echo. Ein Raum mit ernsthaft aufgegriffenem Gegenwind ist offener.

4.2 Operationalisierung

Tabelle 1 zeigt, wie die Dimensionen in messbare Teilindikatoren übersetzt werden. Jede Zeile folgt demselben Muster: eine Leitfrage, ein Label, das pro Kommentar-Einheit vergeben wird, eine offene Formel und eine Schwelle, ab der der Wert als belastbar gilt.

Indikator	Leitfrage	Label je Einheit	Formel	belastbar ab
D1	Wie stark ist der eigene Feed vorsortiert?	Selbstauskunft (4 Fragen)	<i>Mittel der Antworten</i>	2 Antworten
D2.1 Positionale Einigkeit	Ziehen die Stimmen in dieselbe Richtung?	Haltung: zustimmend / widersprechend / neutral	$ 2p - 1 $, $p = \text{Anteil Zustimmender}$	≥ 5 Positionsbezüge
D2.2 Bonding/Bridging	Bleibt der Raum unter sich?	(übernimmt die Labels von D3.4)	$= D3.4$	wie D3.4
D3.1 Quellenframing-Bias	Werden fremde Quellen stärker abgewertet als eigene?	Quelle: eigen / fremd + abgewertet ja/nein	$\text{Anteil abgewerteter Fremdquellen} - \text{Anteil abgewerteter Eigenquellen}$ (min. 0)	≥ 5 Quellen, ≥ 1 fremde
D3.2 Frame-Dichte	Wie dicht sind pauschal abwertende Rahmungen?	Kampfbegriff ja/nein	$\text{Kampfbegriffe} \div \text{wertende Aussagen}$	≥ 10 Aussagen

Indikator	Leitfrage	Label je Einheit	Formel	belastbar ab
D3.3 Ambiguitätsintoleranz	Wird Mehrdeutigkeit ausgehalten?	offen anerkannt / zwanghaft eindeutig / nicht strittig	$zwanghaft \div (offen + zwanghaft)$	≥ 3 strittige Einheiten
D3.4 Bridging (invers)	Werden Gegenpositionen bearbeitet oder abgewiesen?	Gegenposition ja/nein + Rezeption: bearbeitet / abgewiesen / ignoriert	$abgewiesen \div (abgewiesen + bearbeitet)$	≥ 3 Gegenpositionen
D4.1 Gruppenjargon	Ist die Sprache nur für Eingeweihte lesbar?	Jargon ja/nein	$Jargon\text{-}Einheiten \div alle\ Einheiten$	≥ 10 Einheiten
D4.2 Kontextabhängigkeit	Ist die Äußerung ohne Vorwissen verständlich?	kontextabhängig ja/nein	$kontextabhängige \div alle\ Einheiten$	≥ 10 Einheiten

Tab. 1: Operationalisierung. Jede Dimension ist das arithmetische Mittel ihrer belastbaren Teilindikatoren; D3 gilt ab 2 von 4 belastbaren Teilindikatoren, D2 und D4 ab je 1.

Die wichtigste Rechenregel dabei: nie 0, sondern „nicht belastbar“. Fehlt die Datenbasis für einen Indikator, fällt er aus der Rechnung, statt als 0 einzugehen, denn sonst würde Datenmangel als Offenheit fehlgelesen. Ein Raum ganz ohne aufgegriffene Gegenpositionen ist nicht offen, er liefert schlicht kein Signal.

Zwei Erkenntniswege: Vergleich und Referenzdatenbank

Weil es keinen Nullpunkt gibt, gewinnt das Instrument seine Aussagekraft auf zwei Wegen. Erstens durch den direkten Vergleich: Man misst Raum X und Raum Y und sieht für jede Dimension, wo welche Form der Schließung stärker ausgeprägt ist. Zweitens durch eine Referenzdatenbank: Wir haben viele Diskursräume aus unterschiedlichen Bereichen automatisiert erhoben und vermessen, von Politik über Medienkritik und Wissenschaft bis Sport und Unterhaltung, und daraus je Kategorie Mittelwerte und Verteilungen gebildet. Ein neuer Messwert wird in diese Verteilung eingeordnet: unauffällig, solange er unter dem 67. Perzentil des Korpus liegt, erhöht im oberen Drittel, extrem in den obersten zehn Prozent (Abbildung 1). Die Aussage ist immer vergleichend gemeint, also „geschlossener als der typische gemessene Raum“ und nie „absolut geschlossen“.

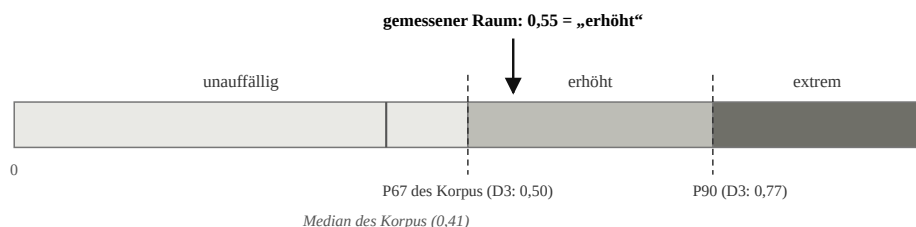


Abb. 1: Einordnung eines Messwerts am Beispiel der Dimension D3. Schwellen und Median stammen aus der laufenden Eichung ($n = 44$ Räume mit belastbarem D3-Wert).

Diese Eichung läuft: Bislang sind 45 YouTube-Kommentarräume quer durch die Kategorien vermessen, und mit jedem weiteren Raum werden die Referenzbereiche schärfer. Die

Kategorien-Profile bleiben dabei ausdrücklich Zwischenstände einer kleinen Stichprobe, keine belastbaren Aussagen über Milieus.

4.3 Funktionsweise

Die Verarbeitungskette verläuft in getrennten Stufen (Abbildung 2): Kommentare über die offizielle Schnittstelle der Plattform laden, die Bezugsposition des Videos erfassen und für den Lauf einfrieren, den Text in Einheiten zerlegen, die Einheiten stapelweise von der KI labeln lassen (ein quelloffenes Großmodell über die hochschuleigene Schnittstelle der UdK), aus den Labels die Werte berechnen und gegen die Referenz einordnen. Jede menschliche Korrektur im Audit fließt als Neuberechnung zurück.

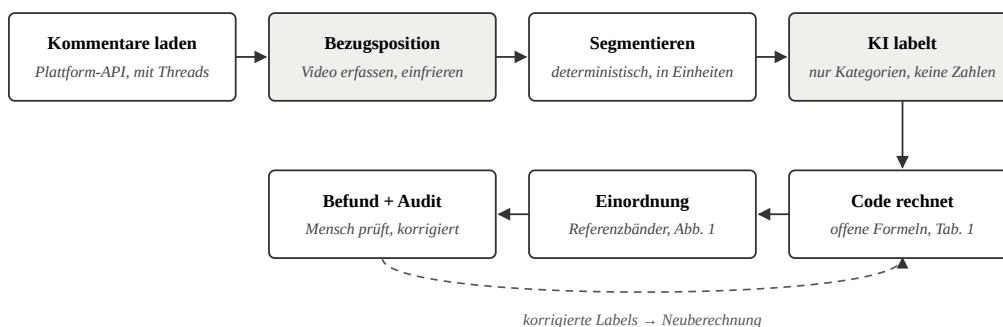


Abb. 2: Die Verarbeitungskette des BSI. Grau hinterlegt sind die KI-Stufen. Die KI vergibt nur Labels, alle Zahlen entstehen deterministisch im Code.

Der Befund selbst hat mehrere Schichten: ein Live-Board mit den wachsenden Dimensionsbalken, ein Meinungsspektrum, in dem jeder Beitrag anklickbar und zum Originalkommentar verlinkt ist, eine automatisch erzeugte, laienverständliche Einordnung der fertigen Werte und ein Vergleichsmodus, der mehrere Räume als Radar übereinanderlegt, inklusive der Referenzprofile aus der laufenden Eichung.

Schwächen der Konstruktion

Auch diese Konstruktion hat Schwächen, die wir offen benennen. Die Kategorien der Referenzdatenbank sind Vereinfachungen; politische Räume links, mitte und rechts sind keine homogenen Milieus, und feinere Unterkategorien wären ein nötiger Ausbau. Mittelwerte glätten die Unterschiede innerhalb einer Kategorie. Schon die Zuordnung beim Erheben ist eine Interpretation, denn viele Inhalte passen in keine Kategorie sauber hinein. Die Stichprobe definiert die Referenz selbst: Ist ein Milieu überrepräsentiert, verschiebt sich, was als normal gilt; die Zusammensetzung des Korpus liegt deshalb offen. Und erhoben werden vor allem die sichtbarsten, von der Plattform nach Relevanz sortierten Kommentare, keine Zufallsauswahl.

5. Was messbar ist, wird (an)greifbar

5.1 Sichtbarmachung statt Gegenrede

Der Titel des Seminars, *Bursting the Bubble*, fragt nach dem Platzen. Dazu zuerst eine Klarstellung: Blasen platzen auch ohne Messinstrument. Gegenrede, Satire oder eine persönliche Begegnung können Räume öffnen, ganz ohne Kennzahl. Der Beitrag eines Messinstruments ist ein anderer. Es macht das Untersuchen und Aufbrechen systematisch. Erst mit einem Maß lassen sich viele Räume vergleichbar vermessen, Interventionen auf ihre Wirkung prüfen und Entwicklungen über Zeit verfolgen. So wie der Gini-Koeffizient Ungleichheit vergleichbar und politisch verhandelbar gemacht hat, soll der BSI Geschlossenheit von der Metapher zur Messgröße machen. Ein nachprüfbarer und kritisierbarer Befund nimmt der gefühlten Geschlossenheit ihre Selbstverständlichkeit, und Diskutierbarkeit ist die Vorbedingung jeder Öffnung.

5.2 Ausblick: Anwendungs- und Forschungspotenzial

Systematische Messbarkeit schafft die Grundlage für eine ganze Familie weiterer Werkzeuge. Am nächsten liegt die ursprüngliche Idee dieses Projekts: aktivistische Kommentar-Tools, die deklarierte Gegenperspektiven in geschlossene Räume einspielen. Der BSI gibt ihnen ein Zielkriterium (wo ist eine Intervention sinnvoll?) und ein Wirkungsmaß (hat sich der Raum danach geöffnet?). Denkbar ist das Instrument auch als Analysewerkzeug für die Plattformen selbst, die ihre Empfehlungs- und Sortierlogik dort anpassen könnten, wo Geschlossenheit extreme Werte erreicht. Moderation geschähe dann nach Befund statt nach Bauchgefühl.

Auch jenseits des Aktivismus eröffnet ein solches Maß Anwendungen. In der politischen Kommunikation ließe sich prüfen, ob Botschaften einen Raum überhaupt erreichen oder dort als Fremdsignal abprallen; die Anschlussfähigkeit der eigenen Sprache würde an konkreten Räumen getestet statt vermutet. Im Wahlkampf könnten Kampagnen die Diskursräume ihrer Zielgruppen kartieren und Ressourcen dort einsetzen, wo Räume noch erreichbar sind. Im Marketing ließe sich beobachten, ob eine eigene Community von lebendiger Einigkeit in abgeschottete Selbstbespiegelung kippt, bevor das als Krise sichtbar wird. Und für die Forschung werden Längsschnittstudien, Plattformvergleiche und die Evaluation von Moderationsmaßnahmen mit einem wiederholbaren Maß erst methodisch sauber möglich.

Zugleich wird dabei eine Doppelverwendbarkeit sichtbar. Dieselbe Kartierung, die Öffnung ermöglichen soll, ließe sich auch nutzen, um geschlossene Räume gezielter zu bespielen. Ein Instrument, das Diskursräume lesbar macht, macht sie für alle lesbar. Umso wichtiger sind die Transparenz des Verfahrens und die offene Dokumentation seiner Grenzen.

6. Reflexion und Fazit

Über die Konstruktionsschwächen aus Kapitel 4.3 hinaus hat der Ansatz Grenzen, die offen benannt gehören. Erstens ist der BSI auf eine Domäne geeicht. Der Kalibrierkorpus und die

Definition der Kampfbegriffe stammen aus dem deutschsprachigen politisch-medialen Diskurs. Auf fachfremde Räume greifen einzelne Indikatoren kaum, sodass diese als unauffällig erscheinen können. Unauffällig heißt dann präzise: wenig von jener Schließung, die das Werkzeug sucht, und gerade nicht offen. Zweitens bleibt D1 eine Selbstauskunft per Fragebogen, weil sich die Kuratierung eines Feeds von außen nicht beobachten lässt. Gemessen werden die Symptome der Vorsortierung, nicht ihre Ursache im Empfehlungssystem. Drittens bleibt jede semantische Etikettierung eine Interpretation. Die Prüfansicht macht sie korrigierbar, aber nicht objektiv. Viertens ist der Referenzkorpus mit 45 Räumen klein und im Aufbau; die Bänder verschieben sich mit jeder Erweiterung, und das ist dokumentiert. Und schließlich ist auch der Maßstab selbst eine Setzung: Dass Geschlossenheit ein Defizit ist und nicht bloß Gemeinschaftsgefühl, folgt nicht aus den Daten, sondern aus einer normativen Haltung, die den Anderen als Gesprächspartner ernst nimmt statt als Störung (Han 2016).

Diese Grenzen relativieren den Anspruch, nicht das Ergebnis. Wir sind angetreten, ein Gegenrede-Werkzeug zu bauen, und haben unterwegs das Problem bearbeitet, das jeder solchen Intervention vorausliegt: eine nachvollziehbare Berechnungsgrundlage dafür, wie geschlossen ein konkreter Diskursraum tatsächlich ist. Das Ergebnis ist kein absolutes Normmaß, sondern ein lauffähiger, nachprüfbarer und kritisierbarer Ansatz, Diskursräume empirisch zu untersuchen. Damit ist die Grundlage für die skizzierten Werkzeuge gelegt, und auch für die Rückkehr zur ursprünglichen Interventionsidee, nun mit einem Maß in der Hand, an dem sich Wirkung überhaupt erst ablesen ließe.

Literatur

1. Foucault, Michel (1991): *Die Ordnung des Diskurses*. Frankfurt a. M.: Fischer Taschenbuch.
2. Habermas, Jürgen (1981): *Theorie des kommunikativen Handelns*. Frankfurt a. M.: Suhrkamp.
3. Han, Byung-Chul (2016): *Die Austreibung des Anderen. Gesellschaft, Wahrnehmung und Kommunikation heute*. Frankfurt a. M.: S. Fischer.
4. Luhmann, Niklas (1984): *Soziale Systeme. Grundriß einer allgemeinen Theorie*. Frankfurt a. M.: Suhrkamp.
5. Mayring, Philipp (2015): *Qualitative Inhaltsanalyse. Grundlagen und Techniken*. 12. Aufl. Weinheim: Beltz.
6. Pariser, Eli (2011): *The Filter Bubble. What the Internet Is Hiding from You*. New York: Penguin Press.
7. Sunstein, Cass R. (2017): *#Republic. Divided Democracy in the Age of Social Media*. Princeton: Princeton University Press.