

Zur Einführung – Datenflut und Informationskanäle

Heike Ortner, Daniel Pfurtscheller, Michaela Rizzolli und Andreas Wiesinger

„Während der Ebbe schrieb ich eine Zeile auf den Sand, in die ich alles legte, was mein Verstand und Geist enthält. Während der Flut kehrte ich zurück, um die Worte zu lesen, und ich fand am Ufer nichts, als meine Unwissenheit.“

(Khalil Gibran)

Anker, Schleusen, Netze – Medien in der Datenflut

Im Digitalzeitalter haben die Produktion, Verbreitung und Speicherung von Daten gigantische Ausmaße angenommen. Pro Minute werden weltweit fast 140 Millionen E-Mails verschickt (Radicati 2014, S. 4), 100 Stunden Videomaterial auf YouTube hochgeladen (YouTube 2014), 350.000 Tweets geschrieben (Twitter 2014), 970 neue Blogeinträge von Wordpress-Usern veröffentlicht (Wordpress 2014) und 240.000 Fotos auf Facebook hochgeladen (Facebook u.a. 2013, S. 6) – Tendenz steigend.

Für die Geistes- und Sozialwissenschaften ergeben sich aus diesen veränderten Kommunikationsverhältnissen neue interessante Forschungsfragen von hoher gesellschaftlicher Relevanz, aber auch methodologische und wissenschaftsethische Probleme: Inwiefern fungieren (welche) Medien in dieser Datenflut (noch) als Gatekeeper? Welche Strategien entwickeln Mediennutzerinnen und -nutzer angesichts der Fülle an Daten? Welche Orientierungsangebote gibt es? Und welche Prozesse müssen Daten durchlaufen, um zu Wissen zu werden? Wer verfügt über welche Daten, wie werden sie weiterverarbeitet und weitergenutzt, wie verschieben sich dadurch Machtgefüge und gesellschaftliche Konzepte von Privatsphäre, Wissen und Verantwortung? Inwiefern können, dürfen und sollen Wissenschaftlerinnen und Wissenschaftler an den produzierten Daten teilhaben, mit welchen Methoden und mit welchen Grenzziehungen werden die Datenberge bewältigt?

Diese Fragen waren der Ausgangspunkt für eine Ringvorlesung zum Thema *Datenflut und Informationskanäle*, die im Wintersemester 2013/14 an der Universität Innsbruck veranstaltet wurde und in deren Rahmen auch der Medientag 2013 mit dem Titel *Anker, Schleusen, Netze – Medien in der Datenflut* stattfand. Organisiert wurden beide Veranstaltungen vom interdisziplinären Forschungsbereich *ims (innsbruck media studies)*. Das erklärte Ziel des Medienforums ist der fachübergreifende Austausch zu aktuellen Brennpunkten der Medienwissenschaft, im Studienjahr 2013/14 eben zu Big Data. Entsprechend wurden in der Ringvorlesung und am Medientag höchst unterschiedliche theoretische Zugänge, kritische Aspekte und Anwendungsbeispiele präsentiert und diskutiert.

Das Thema erwies sich als noch aktueller als ursprünglich gedacht. Zum Zeitpunkt der Konzeption war die globale Überwachungs- und Spionageaffäre, die der ehemalige Geheimdienstmitarbeiter Edward Snowden Ende Juli 2013 enthüllte, noch nicht im vollen Umfang deutlich geworden. In dem Jahr, das seitdem vergangen ist, haben wir sehr viel mehr darüber erfahren, mit welchen Mitteln staatliche Geheimdienste ohne Verdacht Daten über Internetaktivitäten anhäufen und auswerten (einen Überblick gibt Holland 2014).

Vor diesem Hintergrund nimmt sich die vorangegangene gesellschaftliche Verunsicherung angesichts der ‚Macht‘ von Google, Facebook und Co. geradezu gering aus, auch wenn sich das Unbehagen bereits in den Nullerjahren in populären Publikationen niederschlug, die Google und Plattformen des Social Web zur gefürchteten „Weltmacht“ stempelten und vor dem „Google-Imperium“ und der „facebook-Falle“ warnten (vgl. z.B. Reppesgaard 2011, Reischl 2008, Adamek 2011). Dabei sollte mittlerweile allen klar sein, dass es im Internet nichts umsonst gibt, sondern dass man für Gratis-Dienste mit seinen persönlichen Daten bezahlt – überspitzt formuliert: „If you are not paying for it, you're not the customer; you're the product being sold.“ (Lewis 2002)

Selbstaussagen von Google haben nicht unbedingt geholfen, Sorgen dieser Art zu zerstreuen. Eric Schmidt, bis 2001 Geschäftsführer von Google, meinte zu den Zukunftsvisionen seines Konzerns in einem öffentlichen Interview im Rahmen des *Second Annual Washington Ideas Forum* Folgendes:

“With your permission you give us more information about you, about your friends, and we can improve the quality of our searches [...] We don't need you to type at all. We know where you are. We know where you've been. We can more or less know what you're thinking about.” (Thompson 2010)

Angenommen wird von diesen Konzernen und ihren gut finanzierten Forschungsabteilungen, dass man auf der Grundlage des Klick-, Such- und Kommunikationsverhaltens in Google, Facebook, auf Twitter etc. das zukünftige Verhalten von Menschen vorhersagen und in weiterer Folge effektiv steuern kann – welche Produkte jemand kaufen wird, welche politischen Überzeugungen oder sexuellen Vorlieben jemand hat und so weiter. Die Möglichkeiten zur Manipulation und Ausforschung von Konsumentinnen und Konsumenten oder Bürgerinnen und Bürgern scheinen grenzenlos.

Abgesehen von der expliziten Erstellung von Inhaltsdaten sind wir also alle selbst als Medien-nutzende und Konsumentinnen bzw. Konsumenten Datenquellen. Diese Daten sind bereits zu einem monetär relevanten, maßgeblichen Bestandteil gezielten Marketings geworden. Unter dem Schlagwort „Open Data“ wird auch gegenüber dem Staat gefordert, öffentliche Verwaltungsdaten für alle verfügbar und nutzbar zu machen. Gleichzeitig bieten Enthüllungsplattformen à la WikiLeaks gerade geheimen und vertraulichen Daten eine breite Öffentlichkeit. Und auch immer mehr Unternehmen und politische Parteien wollen aus der Datenflut im Netz Profit schlagen. Mit statistisch-algorithmischen Methoden wird beim sogenannten „Data Mining“ versucht, Wissenswertes aus dem Datenberg ans Licht zu befördern.

Als Paradebeispiel für eine Anwendung der Analyse großer Datenmengen wird häufig Google Flu Trends (www.google.org/flutrends/) genannt: Das System zur Voraussage von Grippeerkrankungen basiert auf der Grundlage, dass die Häufigkeit von bestimmten Suchbegriffen mit der Häufigkeit von tatsächlichen Grippefällen zusammenhängt, ein System, das die Häufigkeit von Grippefällen schätzt. Wie die Studie von Lazer u.a. (2014) zeigt, ist der Voraussagealgorithmus bei Weitem nicht perfekt: Zeitweise gab es beträchtliche Abweichungen, im Februar 2013 hat Google Flu Trends vorübergehend mehr als die doppelte Anzahl der tatsächlichen Arztbesuche für grippeähnliche Krankheiten vorausgesagt, während traditionelle statistische Methoden ein viel genaueres Ergebnis erzielten.

Das zeigt auch, welche wichtige Rolle die dahinter stehenden Algorithmen haben. Algorithmen sind für den Erfolg der Unternehmen wichtig und daher nicht offen gelegt, der Zugang zu den Daten selbst ist nicht einfach. Ein beachtlicher Teil der wissenschaftlichen Forschung wird von Mitarbeiterinnen und Mitarbeitern der großen Technologiefirmen Microsoft, Google und Facebook geleistet. Die Forschungsergebnisse selbst sind jedoch publiziert und werden von den Unternehmen öffentlich präsentiert.¹

Algorithmen als Schleusenwärter im Datenstrom

In den Anfangstagen des Kinos waren die sowjetischen Filmregisseure von der Technik der Montage fasziniert. Lew Kuleschow, einer der frühen Theoretiker des Films, vertritt die These, dass das emotionale Potenzial des Kinos nicht in den gezeigten Bildern selbst liegt, sondern in der Art und Weise, wie die Einstellungen zusammenmontiert werden, wobei besonders die Reihenfolge wichtig ist (vgl. Kuleschow 1973). Nach seinen filmischen Experimenten ist der bekannte Kuleschow-Effekt benannt (vgl. Muckenhaupt 1986, 191; Prince & Hensley 1992). Der Kuleschow-Effekt macht uns deutlich, dass das Verständnis des Betrachters bzw. der Betrachterin nicht nur durch den Inhalt des Gezeigten verändert werden kann, sondern auch dadurch, wie diese Inhalte montiert und zusammengeschnitten werden.

Zurück zu Facebook: Man kann die Inhalte des Facebook-Streams mit Einstellungen in einem Film vergleichen. Die einzelnen ‚Einstellungen‘ des Streams (die Postings) zeigen uns weder ‚das wahre Leben‘, noch ist der Stream ‚ungeschnitten‘. Facebook verwendet einen Algorithmus namens EdgeRank, dessen Details nicht bekannt sind, um aus allen Inhalten die für den Nutzer bzw. die Nutzerin relevanten auszuwählen. Der Algorithmus wählt aus, was wir zu sehen bekommen. Vonseiten von Facebook wird betont, das Ziel von EdgeRank sei es, die Relevanz der geposteten Inhalte zu maximieren.

Dass die Auswahl und die ‚Montage‘ der Newsfeeds aber auch andere Ziele haben kann, zeigt das groß angelegte Experiment von Kramer, Guillory & Hancock (2014). Damit konnten sie eine Übertragung von Emotionen (Gefühlsansteckung) auf Facebook nachweisen. Ihre Ergeb-

¹ Bei Facebook unter <https://www.facebook.com/publications>, bei Google unter <http://research.google.com> und bei Microsoft unter <http://research.microsoft.com/>.

nisse lassen vermuten, dass die Emotionen von anderen Userinnen und Usern auf Facebook unsere eigenen Emotionen beeinflussen, was ein Hinweis auf eine Gefühlsansteckung in großen Maßstäben wäre.

Das Experiment wurde in den Medien und von wissenschaftlicher Seite scharf kritisiert. Das wirft auch Fragen zur Ethik von wissenschaftlicher Big-Data-Forschung auf. Big Data bietet der Wissenschaft somit sowohl neue Möglichkeiten als auch neue theoretisch-methodische Herausforderungen. Neben der intensiven Diskussion über den Schutz personenbezogener Daten wird dabei auch die Grundfrage aufgeworfen, was die maschinell-algorithmische Analyse von „Big Data“ ohne Hintergrundwissen über dessen Bedeutung überhaupt leisten kann und wie wir anhand von Daten überhaupt Wissen generieren können.

Als weitere wichtige Schlagwörter im semantischen Netz unseres Themengebietes lassen sich unter anderem Data Mining, Knowledge Discovery in unstrukturierten Daten, Data Warehouse, die Filterblase (*filter bubble*) und Open Data identifizieren. Viele dieser Bereiche werden in den Beiträgen angeschnitten und mehr oder weniger ausführlich behandelt, andere rücken in den Hintergrund. Speziell zum sensiblen Bereich der Privatsphäre und der Offenheit von Daten sei auf einen anderen Sammelband aus der Werkstätte des *ims* verwiesen (vgl. Rußmann, Beinsteiner, Ortner & Hug 2012).

Die Beiträge dieses Bandes im Überblick

Die 13 Beiträge, die in diesem Band versammelt sind, entstanden im Anschluss an die bereits genannte Ringvorlesung im Wintersemester 2013/14 und den Medientag 2013 an der Universität Innsbruck. Zum einen sollen durch diese Publikation die vielen anregenden Vorträge und dadurch angestoßene Denkprozesse dokumentiert werden, zum anderen ist der Sammelband auch ein von diesem Anlass unabhängiger Beitrag zur aktuell brennenden Diskussion um Big Data. Gegliedert ist das Buch in drei thematische Cluster, wobei sich zahlreiche Querverbindungen und gegenseitige Ergänzungen ergeben.

Daten und Netzwerke: Theoretische Verankerungen

Am Anfang steht wie meistens in Sammelbänden dieser Art eine Reihe von Beiträgen, die notwendige Begriffserklärungen vornehmen, das Thema in den weiteren interdisziplinären Forschungshorizont einordnen und methodische Perspektiven aufzeigen.

In seinem Beitrag mit dem Titel *Involuntaristische Mediatisierung. Big Data als Herausforderung einer informativierten Gesellschaft* führt **Marian Adolf** in die Problematik der ‚großen Daten‘ und ihrer Nutzung aus kommunikationswissenschaftlicher Perspektive ein. Als zentrale Fragestellungen arbeitet er die Konsequenzen für die individuelle Privatsphäre und die Informativierung der Lebenswelt auf der Grundlage einer probabilistischen Interpretation von Daten heraus. Er greift den etablierten Terminus Mediatisierung auf und deutet ihn neu: Im Kontext der umfassenden Verwertung von Big Data vor allem zu ökonomischen Zwecken wird die Mediatisierung eine involuntaristische, das heißt ein unbeabsichtigtes Nebenprodukt der

Nutzung nicht nur von Medien, sondern von Services des alltäglichen Lebens (z.B. Kundenkarten). Dass die Analysen, denen jedes Detail unseres Online-Verhaltens unterzogen wird, auch Fehlschlüsse und Missbrauch zulassen, ist eine offensichtliche Problematik, darüber hinaus eröffnen sich aber noch viel fundamentalere Fragen, auf die unsere Gesellschaft noch keine Antwort hat.

Auch **Ramón Reicherts** Beitrag *Big Data: Medienkultur im Umbruch* widmet sich den gesellschaftlichen Auswirkungen der Datenflut, vor allem der vermeintlichen prognostischen Kraft, die der Auswertung von Nutzerverhalten zugesprochen wird. Darüber hinaus werden aber auch grundlegende medienwissenschaftliche Konsequenzen des Social Media Monitoring angesprochen. Konzerne wie Facebook produzieren Metawissen, das zunehmend zur wichtigsten Form von Zukunftswissen wird: Die Wissensbestände beruhen auf Profiling, auf Registrierungs-, Klassifizierungs-, Taxierungs- und Ratingverfahren, die von den Medienunternehmen im Hintergrund (im „Back End“) entwickelt und ausgewertet werden. Die erhobenen Daten sind jedoch bruchstückhaft und von den vorgegebenen Formularstrukturen vorgeprägt, was sich auch auf die Auswertungsergebnisse – wie z.B. den sogenannten Happiness Index von Facebook – auswirkt.

Axel Maireder nimmt in seinem Beitrag *Ein Tweet: Zur Struktur von Netzöffentlichkeit am Beispiel Twitter* einen Tweet zu einer politisch-juristischen Causa (der sogenannten Part-of-the-game-Affäre um Uwe Scheuch) als Ausgangspunkt für eine Diskussion über die Veränderungen unseres Verständnisses von Öffentlichkeit. Der Autor zeigt zudem, wie sich auf Twitter Konversationsgemeinschaften ausbilden, die sich nicht nur durch geteilte politische Einstellungen, sondern durch einen gemeinsamen Kommunikationsstil zusammenfinden. Die verschränkte Mikro- und Makroperspektive macht dabei die hochdynamische Struktur der Netzöffentlichkeit deutlich, die sich erst durch vielfältige und wechselseitige Bezugnahmen unterschiedlicher sozialer Akteure, Gruppen und Konversationen im Diskurs ergibt.

Unterschiedliche Angebote im Web, die vielfach unter dem Begriff „Social Web“ zusammengefasst werden, erzeugen einen scheinbar endlosen Informationsfluss. Der Beitrag *Facebook, Twitter und Co.: Netze in der Datenflut* von **Veronika Gründhammer** widmet sich vor dem Hintergrund eines systemtheoretischen Kommunikationsbegriffs der Frage, inwieweit verschiedene Social-Web-Angebote, die selbst Teil der Datenflut sind, dabei helfen, eben dieser Datenflut Herr zu werden. Dabei wird deutlich, dass Facebook, Twitter und Co. nicht nur einen Strom an Informationen erzeugen, sondern gleichzeitig eine metakommunikative Funktion erfüllen, indem Social-Web-Angebote eine Ordnung und Strukturierung von Medieninhalten begünstigen.

Den Abschluss des theoretischen Teils bilden **Michael Klemm** und **Sascha Michel** mit ihrem Beitrag *Big Data – Big Problems? Zur Kombination qualitativer und quantitativer Methoden bei der Erforschung politischer Social-Media-Kommunikation*. Die beiden Autoren beleuchten vor allem die methodologischen Konsequenzen der Datenexplosion und diskutieren sie anhand des Beispiels ihrer eigenen Erforschung des Verhaltens von Politikerinnen und Politikern in Sozialen Netzwerken wie Twitter. Sie plädieren dafür, die jeweiligen Vor- und Nachteile von

quantitativen und qualitativen Methoden auszunutzen, und stellen Werkzeuge vor, wie man aus der Makroperspektive mit Twitter-Daten umgehen kann, machen aber mit stilpragmatischen Fallanalysen deutlich, dass für bestimmte Twitterstile typische Modalitäten wie ‚Lästern‘ oder ‚Ironie‘ nur interpretativ im Rahmen einer medienlinguistischen Diskurs- und Textanalyse bestimmt werden können.

Bildung und Netzkritik als Schleusen der Informationsflut

In diesem Abschnitt wenden wir uns den Bildungswissenschaften, der Philosophie und der Sprachwissenschaft zu. Gemeinsame Einordnungsinstanz ist die Bedeutung der Datenflut für die Entwicklung des Individuums in Auseinandersetzung mit der realen und der virtuellen Umwelt (wobei diese beiden Sphären ja zunehmend miteinander verschmelzen, wie die vorergehenden Beiträge gezeigt haben).

Ausgehend von der Frage, ob ein Mehr an Daten wirklich einen Mehrwert darstellt, nimmt **Petra Missomelius** in ihrem Beitrag *Medienbildung und Digital Humanities. Die Medienvergessenheit technisierter Geisteswissenschaften* die Verfahren des Data Mining in den Wissenschaften in den Blick. Die Autorin verweist darauf, dass gewonnene Rohdaten erst durch den Menschen, der die Daten auswertet und interpretiert, Sinn erhalten. Wissen lebt von der Wahrnehmung, Bewertung und Verarbeitung durch den Menschen. Im Umgang mit der Datenflut fordert Missomelius eine kritische und produktive Auseinandersetzung mit Instrumenten und Methoden der Digital Humanities, um der Vielfalt, Komplexität und Unschärfe geisteswissenschaftlicher Daten und ihrer technischen Bearbeitbarkeit gerecht zu werden sowie zur Entwicklung fachspezifischer digitaler Werkzeuge und digitaler Forschungsinfrastrukturen beizutragen.

Nach einer grundlegenden Explikation des Datenbegriffs greift **Valentin Dander** in seinem Beitrag *Datendandyismus und Datenbildung. Von einer Rekonstruktion der Begriffe zu Perspektiven sinnvoller Nutzung* ein unscharfes Konzept des Autorenkollektivs Agentur Bilwet auf und deutet es für den Kontext der Medienbildung um. Der Datendandy ist eine Kunstfigur, die unablässig, sprunghaft und ziellos Informationen sammelt und sich dabei von der Datenfülle nicht irritieren lässt. Davon ausgehend beschäftigt sich Dander mit Szenarien der sinnvollen Nutzung von Daten. In diesem Rahmen stellt er die Ergebnisse von Experteninterviews mit NGOs zum Thema vor. Dabei zeigt sich, dass Sammelbegriffe wie ‚Data Literacy‘ mit unklaren Bedeutungszuschreibungen zu kämpfen haben, die vor dem Hintergrund der großen Bedeutung von Datenkompetenz und Datenbildung für die Medienpädagogik kritisch zu sehen sind.

Der Beitrag von **Andreas Beinsteiner** mit dem Titel *Terror der Transparenz? Zu den informationskritischen Ansätzen von Byung-Chul Han und Tiqqun* setzt sich kritisch mit dem Begriff der Transparenz auseinander. Angesichts der wachsenden Datenflut erscheint die Transparenz von Datenerfassung und -auswertung (etwa unter dem Aspekt von Open Data) zwar durchaus sinnvoll, birgt aber auch Risiken, die einer kritischen Reflexion bedürfen. Beinsteiner bezieht sich dabei besonders auf die Positionen des französischen Autorenkollektivs Tiqqun und des Philosophen Byung-Chul Han, die das Dogma der Transparenz hinterfragen.

Heike Ortner beschließt diesen Abschnitt mit dem Beitrag *Zu viel Information? Kognitions-wissenschaftliche und linguistische Aspekte der Datenflut*. Darin geht sie auf kritische Positionen zur Vernetzung und Datenflut ein, darunter auf die Behauptung einer fortschreitenden ‚digitalen Demenz‘ (Manfred Spitzer) und auf den alten Topos des Sprachverfalls. Empirische Ergebnisse aus den Kognitionswissenschaften und aus der Linguistik rechtfertigen weder Jubelstimmung über die von manchen postulierte Steigerung unserer Intelligenz und Leistungsfähigkeit noch eine Verurteilung digitaler Medien als Katalysatoren der allgemeinen Verdummung und mangelhafter sprachlicher Kompetenzen.

Daten in der Praxis

Der letzte Abschnitt dieses Sammelbandes ist Anwendungsbeispielen und empirischen Ergebnissen aus der Big-Data-Forschung gewidmet.

Eva Zangerle behandelt in ihrem Beitrag *Missbrauch und Betrug auf Twitter* Gefahren auf Twitter, der derzeit populärsten Microblogging-Plattform im Netz. Durch seine große Verbreitung wird Twitter zunehmend auch für kriminelle Zwecke missbraucht. Sowohl das Verbreiten von falschen Informationen und von unerlaubter Werbung (sogenanntem „Spam“) wie auch das Hacken von Benutzerkonten werden im Beitrag thematisiert, außerdem stellt Zangerle eine Studie vor, die das Verhalten der Userinnen und User in Fällen der missbräuchlichen Verwendung untersucht.

Mit einer Freizeitbeschäftigung für Jung und Alt, die sich immer größerer Beliebtheit erfreut, befassen sich **Andreas Aschaber** und **Michaela Rizzolli** in ihrem Beitrag *Geocaching – das Spiel mit Geodaten*. Mit der wachsenden Anhängerschaft gehen auch eine steigende Produktion von Daten und ihre materielle Einbindung in Form von Caches in zahlreichen Kultur- und Naturräumen einher. Dabei zeigt sich, dass die hybride Verschränkung virtueller Aktivitäten im physischen Raum zu Konflikten auf unterschiedlichen Ebenen führt. Nicht die Datenflut selbst, sondern der Umgang mit ihr und die Reaktionen darauf sind entscheidend.

Auch für die Translationswissenschaft ergeben sich aus dem massenhaften Anwachsen von Daten Potenziale und spezifische Probleme, wie **Peter Sandrini** in seinem Beitrag *Open Translation Data: Neue Herausforderung oder Ersatz für Sprachkompetenz?* einführend darstellt. Er beschreibt formale und inhaltliche Aspekte von Übersetzungsdaten und ihrer Annotation und die entscheidenden Kriterien für (gute) Open Translation Data. Entsprechende Bemühungen um gesicherte Übersetzungsbausteine können einen wichtigen Beitrag zur Bewältigung von Mehrsprachigkeit leisten.

In seinem Beitrag mit dem Titel *Politische Kommunikation im Social Web – eine Momentaufnahme im Datenstrom* setzt sich **Andreas Wiesinger** mit den Chancen und Risiken des Social Web für die politische Kommunikation auseinander. Plattformen wie Facebook und Twitter werden für die Verbreitung von politischen Botschaften zunehmend wichtiger; neben den Möglichkeiten der digitalen Netzwerke ergeben sich auch Gefahren wie der sogenannte Shitstorm als eine Form der massenhaften Entrüstung im Internet. Wiesinger gibt einen Über-

blick über die verschiedenen Formen der politischen Kommunikation im Social Web und stellt drei Webservices vor, welche die Kommunikate archivieren, auswerten und der Öffentlichkeit zugänglich machen.

Abschließende Bemerkungen und Dank

Bevor wir jenen Stellen Anerkennung zollen, denen sie gebührt, noch eine Anmerkung aus – in Österreich – aktuellem Anlass: Wir als Herausgeberinnen und Herausgeber waren bemüht, eine einheitliche, lesbare und grammatisch korrekte Lösung für das Ansprechen beider Geschlechter zu finden, da wir das grundlegende Anliegen des geschlechtergerechten Formulierens unterstützen. Dabei haben wir fast immer zu Doppelformen gegriffen, da wir davon ausgehen, dass diese nach einer Gewöhnungsphase den Lesefluss ebenso wenig stören wie Literaturhinweise im Text. Wenn es nicht um konkrete Personen, sondern um abstrakte Entitäten geht, und bei einigen wenigen Ausnahmefällen, in denen tatsächlich die Lesbarkeit infrage stand, sind wir jedoch bewusst bei der einfachen Form geblieben.

Wir danken dem Vizerektorat für Forschung, den Fakultäten für Bildungswissenschaften, Mathematik, Informatik und Physik sowie der Philosophisch-Historischen Fakultät und der Philosophisch-Kulturwissenschaftlichen Fakultät der Universität Innsbruck für ihren finanziellen Beitrag zu diesem fakultätsübergreifenden Unterfangen. An der Organisation des Medientages und damit an der Akquirierung und Betreuung einiger unserer Autorinnen und Autoren waren dankenswerterweise Carolin Braunhofer und Magdalena Maier beteiligt. Für ihre sowohl für die Veranstaltung als auch für den Sammelband tragende Unterstützung gilt unseren Kooperationspartnerinnen, der Moser Holding und der Austria Presse Agentur, ganz besonderer Dank. Die bereichernde und reibungslose Zusammenarbeit mit der Moser Holding und mit *innsbruck university press* sowie die Unterstützung durch die Vizerektorin für Forschung sind für die Entstehung dieses Buches von zentraler Bedeutung. Schließlich bedanken wir uns natürlich auch bei den beteiligten Autorinnen und Autoren für ihre Bereitschaft, das Thema ‚Datenflut und Informationskanäle‘ aus so unterschiedlichen Perspektiven aufzugreifen. Den Leserinnen und Lesern wünschen wir, dass dieses Buch wichtige Informationskanäle öffnet und dass einige Gedanken als Rettungsanker in der Datenflut dienen können.

Literatur

Adamek, Sascha (2011): *Die facebook-Falle. Wie das soziale Netzwerk unser Leben verkauft*. München: Heyne.

Facebook, Ericsson & Qualcomm (2013): *A Focus on Efficiency. A whitepaper*. Abgerufen unter: <http://internet.org/efficiencypaper> [Stand vom 30-07-2014].

Holland, Martin (2014): *Was bisher geschah: Der NSA-Skandal im Jahr 1 nach Snowden*. <http://www.heise.de/newsticker/meldung/Was-bisher-geschah-Der-NSA-Skandal-im-Jahr-1-nach-Snowden-2214943.html> [Stand vom 30-07-2014].

- Kramer, Adam D. I.; Guillory, Jamie E. & Hancock, Jeffrey T. (2014): Experimental evidence of massive-scale emotional contagion through social networks. *Proceedings of the National Academy of Sciences* 111 (24), S. 8788–8790.
- Kuleshow, Lew (1973). The Origins of Montage. In: Schnitzer, Luda und Jean et al. (Hrsg.): *Cinema in Revolution: The Heroic Era of the Soviet Film*. London: Secker & Warburg, S. 65–76.
- Lazer, David; Kennedy, Ryan; King, Gary & Vespignani, Allesandro (2014): The Parable of Google Flu: Traps in Big Data Analysis. *Science* 343 (6176) (March 14), S. 1203–1205.
- Lewis, Andrew (2002): *Kommentar auf Metafilter*. Abgerufen unter: <http://www.metafilter.com/95152/Userdriven-discontent#3256046>. [Stand vom 30-07-2014].
- Muckenhaupt, Manfred (1986): *Text und Bild. Grundfragen der Beschreibung von Text-Bild-Kommunikationen aus sprachwissenschaftlicher Sicht*. Tübingen: Narr.
- Prince, Stephen & Hensley, Wayne E. (1992): The Kuleshov Effect: Recreating the Classic Experiment. *Cinema Journal* 31 (2), S. 59–75.
- Radicati (2014): *Email Statistics Report, 2014-2018*. Abgerufen unter: <http://www.radicati.com/wp/wp-content/uploads/2014/04/Email-Statistics-Report-2014-2018-Executive-Summary.docx> [Stand vom 30-07-2014].
- Reppesgaard, Lars (2008): *Das Google-Imperium*. Hamburg: Murmann.
- Reischl, Gerald (2008): *Die Google-Falle. Die unkontrollierte Weltmacht im Internet*. Wien: Ueberreuter.
- Rußmann, Uta; Beinstainer, Andreas; Ortner, Heike & Hug, Theo (Hrsg.) (2012): *Grenzenlose Enthüllungen? Medien zwischen Öffnung und Schließung*. Innsbruck: iup (= Edited Volume Series).
- Thompson, Derek (2010): Google's CEO: "The Laws Are Written by Lobbyists". *The Atlantic*. Abgerufen unter: <http://www.theatlantic.com/technology/archive/2010/10/googles-ceo-the-laws-are-written-by-lobbyists/63908/> [Stand vom 30-07-2014].
- Twitter (2014): About Twitter, Inc. Abgerufen unter: <https://about.twitter.com/company> [Stand vom 30-07-2014].
- Weinberger, David (2013): *Too big to know*. Aus dem amerikanischen Englisch von Jürgen Neubauer. Bern: Huber.
- Wordpress (2014): *Stats*. Abgerufen unter: <http://wordpress.com/stats/> [Stand vom 30-07-2014].
- YouTube (2014): *Statistics*. Abgerufen unter: <https://www.youtube.com/yt/press/statistics.html> [Stand vom 30-07-2014].