

Phänomeno-Logik

Mediale Praktiken mit Datenbanken

Galt das vorangegangene Kapitel der medienwissenschaftlichen Analyse technischer Aspekte von Datenbanken, sollen im Folgenden mediale Praktiken mit Datenbanken im engen und weiten Sinn betrachtet werden. Das vordergründige Ziel ist hierbei nicht die Auseinandersetzung mit konkreten Datenbanken, sondern die phänomenologische Beschreibung von möglichen Nutzungsweisen von Datenbanken. Untersucht werden soll, wie sich digitale Datenbanktechnologien in mediale Praktiken einschreiben und wie sie diese strukturieren. Ein Anspruch auf Vollständigkeit wird hierbei aufgrund der Vielgestaltigkeit unterschiedlicher Datenbankpraktiken in der digitalen Medienkultur nicht erhoben, aber es werden zentrale Dimensionen des Datenbankgebrauchs freigelegt, an denen sich nicht zuletzt auch die konstatierte Variabilität der Gebrauchs- und Erscheinungsformen digitaler Datenbanken zeigt.¹

1 | Die phänomenologische Betrachtung basiert auf Einklammerung und Variation. Sie stellt die Phänomene in der Epoché still, um diese im »Wie ihres Erscheinens zu fassen« (Günzel 2009: 158) oder dem »Wie des Auftretens« (Günzel 2009: 159); d.h. der phänomenologisch Forschende interessiert sich nicht dafür, was etwas ist oder was es bedeutet, sondern dafür, wie es erscheint oder auftritt, und dafür, »dass es Bedeutung hat« (Günzel 2009: 160). Vor dem Hintergrund des *linguistic turn*, aber auch (post-)strukturalistischer und konstruktivistisch-systemtheoretischer Theorieentwürfe wurde die Phänomenologie häufig für ihre bewusstseinsphilosophische Grundausrichtung sowie ihre Subjektzentrierung kritisiert, wie sie in der Begründung der Phänomenologie bei Edmund Husserl zentral angelegt ist. Dieses Manko haftet der Phänomenologie heute noch immer an, weshalb diese häufig allzu schnell abgetan wird und fruchtbare Anschlussmöglichkeiten sowie Kontinuitäten übersehen werden (vgl. Ihde 2003b: 8f.). So zeigt Stephan Günzel, dass Foucaults Methode der Diskursanalyse durchaus Parallelen zum phänomenologischen Vorgehen aufweist: »Ebenso wie Foucault einen Text nicht auf seine Referenz hin liest, sondern die Weise seines Auftretens analysiert, beschreiben Phänomenologen Dinge oder Objekte unter Absehung der Frage, was sie

Die Herausforderung einer Beschreibung der Phänomeno-Logik digitaler Datenbanken kann in Anlehnung an Heideggers in *Sein und Zeit* entwickeltes Verständnis der Phänomenologie auf die Frage gebracht werden: Wie zeigt sich die Datenbank je nach Zugangsart zu ihr von sich selbst her?² Diese Frage weist in Richtung des relationalen Verhältnisses zwischen der Datenbank als einer medialen Konstellation bzw. als Sammlung medialer Konstellationen und ihrem Gebrauch in medialen Praktiken durch die Nutzer. Im Anschluss an Don Ihdes postphänomenologische Studien lässt sich dieses Verhältnis als ein zweifaches *Embodiment* beschreiben, welches er am Beispiel der instrumentellen Wahrnehmung in den Wissenschaften beschreibt: »[S]cience is necessarily ›embodied‹ in technologies or instruments, but simultaneously it implicates human embodiment as that to which the ›data‹ are reflected« (Ihde 2003a: 133).³ Analog hierzu können Benutzerschnittstellen als Formen der Verkörperung von Datenbanken begriffen

jeweils für jemanden bedeuten mögen« (Günzel 2009: 157). Wovon sich Foucault distanziert ist demzufolge das Begründungsunterfangen Husserls, aber nicht das phänomenologische Vorgehen: »Diskursanalyse ist«, so die These Günzels, »Phänomenologie minus deren Begründung« (Günzel 2009: 160).

2 | Bei seiner Bestimmung des phänomenologischen Phänomenbegriffs rekurriert Heidegger auf die antike Philosophie. Phänomene seien »die Gesamtheit dessen, was am Tage liegt oder ans Licht gebracht werden kann« (Heidegger 1993 [1927]: 28). Die so bestimmten Phänomene wurden auch als das Seiende bezeichnet, welches sich »in verschiedener Weise, je nach Zugangsart zu ihm, von ihm her selbst zeigen« (Heidegger 1993 [1927]: 28) kann. Um eine Beschreibung der Weise des Sich-Zeigens ist die Phänomenologie bemüht, wie auch im Übergang vom vulgären und phänomenologischen Phänomenbegriff deutlich wird: Bedeutet Phänomen im Normalfall »Sich-an-ihm-selbst-zeigen«, so richtet sich die Phänomenologie auf das »Sich-so-an-ihm-selbst-zeigende« (Heidegger 1993 [1927]: 31), d.h. auf Anschauungsformen.

3 | Es ist anzumerken, dass Ihde das Unbehagen vieler Kritiker an der Setzung des intentionalen Bewusstseins sowie des transzendentalen Egos als Ausgangs- und Zielpunkt phänomenologischer Beschreibungen teilt. Dennoch hält er an der phänomenologischen Methode fest, die er strategisch als *Postphänomenologie* bezeichnet, um die Differenzen zu der am Bewusstsein und Subjekt orientierten Phänomenologie husserlscher Prägung zu unterstreichen. An die Stelle des Bewusstseins setzt Ihde im Anschluss an Merleau-Ponty das Embodiment, an die Stelle des Subjekts die Erfahrung und an die Stelle der Wesensschau die auf empirischen Erfahrungen gründende Beschreibung multistabiler Phänomene (vgl. Ihde 2008: 6f.). Zentrales Moment ist für ihn die Relationalität und Reflexivität der Welterfahrung (vgl. Ihde 2003a: 133). Insbesondere in den Wissenschaften zeige sich, dass die Erfahrung der Welt nicht unmittelbar ist, sondern vermittelt. Welterfahrung ist heute zunehmend eine Erfahrung durch Instrumente, in denen die Wissenschaften verkörpert sind.

werden, die auf den Erlebnis- und Erfahrungshorizont ihrer Nutzer (seien es menschliche oder technische Akteure) hin entworfen sind und ihnen im gleichen Zuge eine bestimmte Position zuweisen, aus der heraus sie die Datenbank erfahren und mit dieser umgehen können. Es handelt sich um ein doppeltes Bedingungsverhältnis, bei dem die Verkörperung der Datenbank einerseits davon abhängt, wie ihre antizipierten Nutzer erfahren und handeln können. Andererseits werden die Möglichkeiten der Nutzer, mit der Datenbank umzugehen, aber auch vom Interface bedingt respektive bestimmt. Für die Analyse der Phänomeno-Logik digitaler Datenbanken stellt die Betrachtung von Schnittstellen infolgedessen einen zentralen Referenzpunkt dar. Zugleich geht die verfolgte Fragestellung aber auch über eine beschreibende Analyse von partikularen Benutzerschnittstellen hinaus, da elementare Gebrauchsweisen von Datenbanken in den Blick genommen werden sollen, welche in unterschiedlichen Schnittstellen verkörpert werden können.

Dem Gebrauch von Datenbanken und von den in ihnen versammelten Informationen wird sich im Folgenden auf drei Ebenen zugewandt. Zunächst wird die Operationsweise der Datenbank als latente Infrastruktur betrachtet. Im zweiten Teil wird das Finden des Einen im Vielen thematisiert, bevor im abschließenden dritten Teil die Auswertung des Vielen beleuchtet wird, d.h. die Analyse von Informationssammlungen mit dem Ziel, neue Informationen zu entdecken.

Fungieren Datenbanken als latente Infrastrukturen, bleiben diese als Datenbank für die Nutzer zumeist unsichtbar. Jedoch schreiben sie sich aus der verborgenen Tiefe des Computers heraus an der Oberfläche ein und beeinflussen die Weisen des Umgangs mit digitalen Medienobjekten. Ihre mediale Logik zeigt sich mittelbar in den medialen Praktiken mit digitalen Informationen, welche durch Datenbanktechnologien ermöglicht werden. Bei der Suche nach Informationen in Informationssammlungen und deren Auswertung zeigt sich die Datenbank demgegenüber *als* Datenbank, aus der Informationen geschöpft werden können. Die Benutzerschnittstelle dient als Projektionsfläche, auf der die Inhalte aus der Datenbank zur Erscheinung kommen und auf der die Informationspotenziale der Datenbank inszeniert werden. Sie strukturiert den Zugriff auf den Informationsbestand und bedingt den Umgang mit diesen Informationen.

Unabhängig davon, ob Datenbanken in ihrem Gebrauch *als* Datenbank zur Erscheinung kommen oder nicht, lassen sich diese Gebrauchsformen digitaler Datenbanken im Rahmen der von Luhmann vorgeschlagenen medialen Topologie des Computers beschreiben. Datenbankinformationen erweisen sich als relativ autonome (individuell adressierbare und vielfältig verarbeitbare) Entitäten, die losgelöst von ihrer Erscheinungsform auf der Oberfläche in der Tiefe des Computers Bestand haben. In Anlehnung an Heideggers Technikphilosophie können Datenbanken demzufolge als Technologien beschrieben werden, die Informationen in einen Informationsbestand transformieren (vgl. Heidegger 2000 [1953]: 20f.). Dieser Bestand ist jedoch nur ein Potenzial, welches in verschiedenen medialen Praktiken aktualisiert werden muss. Obwohl bereits In-Formation, dienen Datenbankinformationen als quasi-formloses Medium für Formbildung an der Oberfläche des

Computers.⁴ In diesem Sinne oszillieren Datenbanken unentschieden zwischen medialer Konstellation – oder genauer: Sammlung medialer Konstellationen – und medialer Konfiguration. Denn Datenbanktechnologien eröffnen einen dreifachen Möglichkeitsraum, nämlich die variable Präsentation, Selektion und Auswertung von gespeicherten Informationen. Durch digitale Datenbanken werden Informationen zu Medien für Formbildungen, welche sich an und durch Benutzerschnittstellen ausformen. Hierin liegt die funktionelle Offenheit von Datenbanken begründet, welche als zentrales Gestaltungsprinzip bereits in den Entwurf der ANSI/X3/SPARC-Datenbankarchitektur eingegangen ist. Das Ausgangsproblem, dem sich die *Study Group on Data Base Systems* zuwandte, war die Vermittlung zwischen der internen Speicherlogik des Computers (oder Speicherlogiken verschiedener Computersysteme) und den vielfältigen Gebrauchslogiken ihrer Nutzer. Deswegen wurden Datenbanktechnologien von Anfang an daraufhin entworfen, nicht nur ein Benutzerinterface, sondern viele Interfaces zu haben, die einen je eigenen Blick auf die Datenbank eröffnen und unterschiedliche Formen des Umgangs mit den verfügbaren Informationen ermöglichen sollen. Technisch ist die Möglichkeit zur Vervielfachung von Schnittstellen das Resultat der zunehmenden Entkopplung von Informationen und Programmen, d.h. der technisch gestützten Autonomisierung digitaler Informationen:

»In the separation of data from programs, information becomes a quantity that is not only universal (it can be used by several programs) but also polyvalent (it can be put to several uses), data become morphological, capable of being instituted in multiple ways.« (Thacker 2006: 108)

Die Weise der Präsentation von Datenbankinformationen sowie die Verkörperung von Datenbanken *als* Datenbanken an der Benutzeroberfläche sind optional, d.h. Datenbanken und die in ihnen gespeicherten Informationen können stets auch anders zur Erscheinung kommen. Das Interface tritt zu den Informationen hinzu, gibt diesen eine Form und strukturiert den Umgang mit ihnen. Den Informationen aus der Datenbank wird durch die Benutzerschnittstelle eine Form aufgepfropft.⁵ Zur Erscheinung kommen – paradox formuliert – informierte Informationen. In dieser Hinsicht erweist sich das im Kapitel »Datenbankmodelle« (S. 221ff.) herangezogene Konzept der symbolischen Form als problematisch, um mediale Praktiken mit Datenbanken zu beschreiben. Cassirers Betrachtung symbolischer Formen zielt darauf ab, historisch wandelbare, aber relativ stabile Erkenntnisformen freizulegen

4 | Diese Formulierung ist an Luhmanns Medium/Form-Unterscheidung angelehnt, welche im Kapitel »Medium« (S. 62ff.) ausführlich diskutiert wurde. Jedoch können Datenbankinformationen zugleich als Medium und als Form begriffen werden – eine Möglichkeit, die Luhmann eigentlich nur dem allgemeinen Medium Sinn zuschreibt (vgl. Luhmann 1995: 174).

5 | Zur Logik der Pfropfung siehe Wirth (2006a, 2011b, a).

und zu analysieren. Die durch das konzeptuelle Schema oder durch algorithmische Auswertungsverfahren symbolisch geformten Informationen sind jedoch allenfalls latente Informationen, die in der medialen Praxis symbolisch umgeformt, überformt, wenn nicht sogar verformt werden können. Diese *Transformierbarkeit* bzw. *Transformation* bereits geformter Informationen gilt es medientheoretisch in den Blick zu nehmen. Die auf die Unterscheidung von kulturell bedingten Formen der Erkenntnis abzielende Philosophie symbolischer Formen stellt hierfür kein Vokabular bereit.

Nutzer werden in der medialen Praxis mit Datenbanken nicht in einem Informationsraum situiert, sondern gegenüber einem kontingenten Interface platziert, welches Datenbankinformationen auf eine gewisse Weise zur Darstellung bringt und den Umgang mit diesen Informationen auf unterschiedliche Weise bedingt. Dem gilt es im Folgenden nachzugehen, indem die Vielgestaltigkeit und der Facettenreichtum dieser medialen Praktiken aufgezeigt wird. Es geht hierbei um eine analytische Beschreibung von zentralen Aspekten medienkultureller Entwicklungen sowie von Problemen, die mit der Datenbank als latenter Infrastruktur, dem Finden des Einen im Vielen sowie der Auswertung des Vielen einhergehen. Den größten Raum nimmt dabei der dritte Teil ein, da den Auswertungsmöglichkeiten von Informationssammlungen in der aktuellen Debatte um Big Data große Aufmerksamkeit geschenkt wird. Dennoch muss zugunsten der folgenden Überblicksdarstellung eine detaillierte Analyse ausbleiben. In der Zusammenschau dieser heterogenen Gebrauchsweisen wird jedoch deutlich, dass Datenbanken jeden Versuch der Zuschreibung einer einheitlichen Logik unterlaufen.

DIE DATENBANK ALS LATENTE INFRASTRUKTUR

Der Default-Modus der Versammlung *von* respektive des Umgangs *mit* Informationen in Datenbanken ist das Formular, welches als Eingabe-, Such- und Ergebnismaske auf der Oberfläche erscheint.⁶ Das Formular dupliziert die Struktur von Datenbankeinträgen, für deren serialisierte Speicherung und Verarbeitung die Tabelle im relationalen Datenmodell als »conceptual representation« (Codd 1982: 111) dient. Es bringt die diskrete, diskontinuierliche, segmentierte Einheit von Datensätzen zur Darstellung, welche Deleuze in *Postskriptum über die Kontrollgesellschaften* als »dividuelle Einheit« charakterisiert hat: »Die Individuen sind ›dividuell‹ geworden, und die Massen Stichproben, Daten, Märkte oder ›Banken‹« (Deleuze 1993: 258).

Die Felder eines Formulars bilden adressierbare Einheiten, deren Bedeutung oder Funktion formal durch die Benennung des Felds, d.h. durch ihre Legende bestimmt wird. In der Darstellungsform des Formulars kommen damit die Prinzipien

6 | Kulturhistorisch betrachtet dienten Formulare vor allem als Verwaltungs- und Rationalisierungsinstrumente; siehe hierzu exemplarisch Becker (2008).

der computertechnischen Verwaltung von Informationen in Datenbanken auf paradigmatische Weise zur Erscheinung: die logisch-formale Segmentierung von Informationen und die attributive Beschreibung von Entitäten. Der durch Formulare an der Oberfläche realisierte Darstellungs- und Interaktionsmodus ist analog zu (aber nicht notwendig identisch mit) der Form der Strukturierung und Verwaltung der Information in der Tiefe der Datenbank. Infolgedessen ist es wenig überraschend, dass Formulare nicht nur die medienhistorisch erste, sondern auch eine noch immer dominante Schnittstelle für die Eingabe, Abfrage und Ausgabe von Datenbankinformationen am Bildschirm sind.⁷

Indem Formulare Felder vorgeben, die von den Nutzern ausgefüllt werden können bzw. müssen, strukturieren diese nicht nur die Interaktion mit Datenbanken, sondern entfalten eine Form, die die Datenbank supplementiert, indem ihr das Formular an der Oberfläche eine spezifische (kommunikative) Funktion aufpfropft, welche die Datenbank als Tiefenstruktur überdeckt und mitunter sogar verbirgt.⁸ Deutlich zeigt sich dies, wie Ramón Reichert in *Amateurs im Netz* analysiert hat, an den Profilen auf sozialen Netzwerkseiten wie Facebook, Google+, LinkedIn oder Xing, deren Formularstruktur die nutzerseitige Selbstpräsentation unter Regeln stellt, die bedingen, was im Rahmen eines Nutzerprofils auf welche Weise gesagt und getan werden kann (vgl. Reichert 2008: 95ff.). Obwohl sich ein Nutzer durch die Registrierung in einem sozialen Netzwerk in eine Datenbank einschreibt und zu einem Datenbankeintrag wird, entwirft er an der Oberfläche ein Nutzerprofil und arbeitet an der kommunikativen Vernetzung mit anderen Nutzern. Daher dient nicht die Datenbank, sondern das Netzwerk als leitende Gestaltungsmetapher derartiger Anwendungen. Das Netzwerken als Form medialen Handelns

7 | Tabelle und Formular geben Daten eine Form und verwandeln diese hierdurch »in In-Formationen [...], die eine automatische Verarbeitung ermöglichen« (Krajewski 2007: 39). Medienhistorische Vorläufer dieser Formalisierung sind Krajewski zufolge Katalogkarten, Karteikarten und Lochkarten, die Daten in Form bringen (informieren), indem sie diese in ein »*Diagramm* mit vorgefertigten Positionen [eintragen, M.B.], an denen die Informationen maschinell abzulesen (oder besser: abzutasten) sind« (Krajewski 2007: 45). Die feste Positionierung von Informationen wird in digitalen Medien jedoch optional, d.h. Form und Material sind nicht mehr notwendig in Deckung zu bringen, um eine automatische Verwaltung und Verarbeitung zu gewährleisten. An die Stelle der festen Positionierung tritt die assoziative Adressierung von Informationen durch die Angabe von Spalten und Zeilen, die infolgedessen dynamisch umgeordnet werden können.

8 | Die Felder von Formularen sind vorgegebene *slots*, die durch den Nutzer ausgefüllt werden können bzw. müssen (*filler*): »Formulare bestehen aus einem in vielfach reproduzierter Version vorliegenden Schriftsatz mit Aufforderungen zu bestimmten schriftlichen Handlungen, die in Form und Inhalt eng festgelegt sind (nach dem Muster *slot and filler*)« (Weingarten 1994: 159f.). Zur *slot*- und *filler*-Funktion von Formularen siehe auch Reichert (2008: 94ff.).

und als Imperativ überlagert die Datenbank, welche als latente Infrastruktur in den Hintergrund tritt.

Obwohl Formulare als eine zentrale Form der Präsentation *von* und des Umgangs *mit* Datenbankinformation betrachtet werden können, eröffnen diese keinen unvermittelten oder sogar privilegierten Blick auf die Datenbank in der unsichtbaren Tiefe des Computers. Zwar wiederholen Formulare die abstrakte Logik der Segmentierung von Informationen und der Attribuierung von Bedeutung an der Benutzeroberfläche, doch sie können die logische und physische Gesamtstruktur des Informationsbestands in der Tiefe des Computers auch maskieren. Die Präsentation von Datenbankinformationen als Formular basiert eben auch auf einer Übersetzung zwischen Oberfläche und Tiefe, die zumeist nicht nur selektiv ist, sondern auch etwas auf dem Bildschirm zur Erscheinung bringen kann, was so nicht in der Datenbank gespeichert ist. Gibt die Anfrage an eine Personaldatenbank beispielsweise das Alter eines Angestellten zurück, dann ist davon auszugehen, dass in der Datenbank nicht eigentlich das Alter, sondern das Geburtsdatum einer Person gespeichert ist, aus dem das momentane Alter abgeleitet und an der Oberfläche zur Darstellung gebracht werden kann. Nur weil Datenbankinhalte an der Oberfläche in Form eines Formulars (oder einer Tabelle) verkörpert werden, erhält man keinen unverstellten Blick auf die Datenbank. Einem Formular allein ist nicht anzusehen, wie genau die Informationen in der Tiefe des Computers strukturiert sind. Es handelt sich um eine mögliche Form, die Datenbankinformationen an der Benutzeroberfläche gegeben werden kann. In dieser Variabilität besteht ein zentrales Charakteristikum digitaler Datenbanken. Sie machen die in ihnen versammelten Informationen auf unterschiedliche Weise nach außen anschlussfähig. Dementsprechend basiert vieles, was Nutzer *in* und *mit* digitalen Medien im Allgemeinen und im Internet im Besonderen tun und erfahren, auf digitalen Datenbanktechnologien, ohne dass diese sich *als* Datenbank zeigen. Jeder Blog und die meisten zeitgenössischen Webseiten ebenso wie Web 2.0-Services beruhen auf Datenbanken, die im Hintergrund operieren. Während die Architektur des WWW es ermöglicht hat, Dokumente in einem globalen hypertextuellen Dokumentnetzwerk zu adressieren, machen Datenbanken Informationen *als* Informationen adressierbar.

Löst der Hypertext die Einheit des Texts auf, dann lassen sich Datenbanken als Radikalisierung von Hypertexten verstehen, da sie zudem mit der Einheit des Dokuments brechen.⁹ Obwohl Hypertexte den *einen* Text in *viele* Texte segmen-

9 | Unter den Bedingungen digitaler Medientechnologien wird der Begriff des Dokuments, wie Buckland festgestellt hat, in zunehmenden Maße problematisch. Im Strom digitaler Daten geht die materielle Einheit von Dokumenten verloren, weshalb es einer erneuten Reflexion darüber bedarf, was ein Dokument ist: »A conventional document, such as a mail message or a technical report, exist physically in digital technology as a string of bits, but so does everything else in a digital environment« (Buckland 1997: 808). Zur Klärung der Frage schlägt er die Betrachtung der weithin vergessenen Debatte um den Begriff des Dokuments vor, welche in der ersten Hälfte

tieren, bleiben sie der formalen Materialität des Dokuments verhaftet.¹⁰ Der Logik des Dokuments stellt die Datenbank die diskontinuierliche und diskrete Vielheit von Datenbankeinträgen gegenüber. In dieser Vielheit besteht die Einheit des Dokuments allenfalls als unbestimmter und nicht formalisierbarer Rest in Form von Binary Large Objects, kurz: BLOBs, fort. An der Oberfläche erweist sich das Dokument jedoch als ein möglicher Darstellungsmodus, welcher die Informationspartikel aus der Datenbank in die formale Einheit eines Dokuments zusammenzieht respektive sie als solche präsentiert.

Die Adressierbarkeit von digitalen Informationen eröffnet Alan Liu zufolge einen »encoded or structured discourse« (2008: 209), für den die Transformierbarkeit von Informationen, ihre autonome Mobilität und automatische Verarbeitbarkeit charakteristisch ist (vgl. Liu 2008: 216). Auf Grundlage der strukturierten Speicherung von Informationen in der Tiefe des Computers können Inhalt und Form an der Oberfläche weitgehend losgelöst voneinander verändert werden. Auch wenn sich die Datenbank als latente Infrastruktur mitunter selbst nicht zeigt, bedingt sie die möglichen Weisen des Umgangs mit Informationen an der Oberfläche. Die phänomenale Logik der Datenbank zeigt sich hierbei mittelbar im Gebrauch oder Umgang mit digitalen Informationen, wenn z.B. Inhalte auf verschiedenen Geräten (Laptop, Tablet-Computer, Smartphone etc.) unterschiedlich präsentiert werden oder die Nutzer sich per Knopfdruck dieselben Informationen auf unterschiedliche Weise anzeigen lassen können. Als technische Infrastruktur bedingt sie die Artikulation, Manipulation, Präsentation und Zirkulation von medialen Konstellationen als Informationen in und mit Computern.

Für Alan Liu ist das Streben nach einer Trennung von Inhalt und Form in digitalen Medien ein zentraler Bestandteil des von ihm so bezeichneten *Auf-*

des 20. Jahrhunderts im Kontext der Entwicklung der Dokumentationswissenschaft geführt wurde. In Rekurs auf Paul Otlet, Suzanne Briet und Walter Schürmeyer zeichnet Buckland das Aufkommen eines funktionellen Dokumentbegriffs nach, demzufolge etwas – ein Ding – zu einem Dokument wird, wenn es als solches fungiert. Als charakteristische Funktion von Dokumenten wurde ihr Belegcharakter erachtet, wie Buckland vor allem in Rekurs auf Briet darlegt. Bei der Frage, ob etwas ein Dokument konstituiert, sind vier Kriterien zu bedenken: Dokumente sind Dinge (Materialität), die als Beleg oder Aussage für etwas dienen sollen (Intentionalität) und die infolgedessen in ein Dokument verwandelt wurden (Prozessierung), welches in letzter Instanz auch als solches wahrgenommen werden muss (Phänomenologie) (vgl. Buckland 1997: 806).

10 | Kirschenbaum unterscheidet die formale Materialität von der forensischen Materialität digitaler Medien. Letztere richtet sich auf das Material im engeren Sinn, auf die Hardware und die physischen Inskriptionen digitaler Informationen auf Datenträgern. Als formale Materialität bezeichnet er demgegenüber die Eigenlogik, die medialen Konstellationen durch Software, Formate etc. eingeschrieben wird (vgl. Kirschenbaum 2008: 132ff.).

schreibesystems 2000.¹¹ Digitale Datenbanken und die in ihrer Architektur angelegte Loslösung der computerseitigen Speicherordnung von den nutzerseitigen Gebrauchsordnungen exemplifizieren die Trennung von Inhalt und Form auf paradigmatische Weise. Explizit formuliert wurde der Wunsch, die Erscheinungsweise digitaler Informationen unabhängig, d.h. in Absehung des konkreten Inhalts, verändern zu können, jedoch zunächst in einem anderen Kontext (vgl. Liu 2008: 216). Parallel zur Entwicklung von Datenbankmanagementsystemen setzte Ende der 1960er Jahre die Entwicklung von Markupsprachen ein, die den Übergang von einem prozeduralen zu einem deskriptiven Paradigma der Textverarbeitung markiert, wie Charles Goldfarb – einer der Väter der *Standardized General Markup Language* (SGML)¹² – darlegt (vgl. Goldfarb 1981).¹³ In seinem kurzen Abriss der Geschichte von Markupsprachen schreibt Goldfarb die Idee der Trennung von Inhalt und Form digitaler Texte William Tunncliffe zu, der diese Idee erstmals 1967 im Rahmen eines Vortrags formuliert haben soll (vgl. Goldfarb 1995: 567f.).¹⁴ Die dominante Markupsprache ist heute XML (eXtensible Markup Language), welche unter anderem auch die Basis von HTML (HyperText Markup Language) bildet.

Markupsprachen ermöglichen die Trennung von Inhalts- und Präsentationsform durch die logisch-deskriptive Auszeichnung von Dokumenten. Hierbei werden die Dokumentinhalte um computerlesbare (oder genauer computerverarbeitbare) Metainformationen ergänzt, indem beispielsweise Überschriften von einem Überschrift-Element `<headline> ... </headline>` oder Textabsätze von einem Paragraphen-Element `<paragraph> ... </paragraph>` umschlossen werden (vgl. Lobin 2000: 9ff., 37ff.).¹⁵ Wie die derart ausgewiesenen Elemente eines Dokuments

11 | Liu schließt damit an Kittlers Beschreibung der Aufschreibesysteme 1800/1900 an, wobei seines Erachtens die Trennung von Inhalt und Form die leitende Ideologie des aktuellen Aufschreibesystems darstellt (vgl. Liu 2008: 211f.).

12 | Die Auszeichnungssprache SGML wurde 1986 zu einem offiziellen ISO-Standard. Aus dieser ist Ende der 1990er Jahre die *eXtensible Markup Language* (XML) hervorgegangen (vgl. Lobin 2000: 2f.).

13 | Wie Liu in Bezug auf die USA darlegt, wurde die Nutzbarmachung von Markupsprachen in den Geistes- und Kulturwissenschaften zu dieser Zeit vor allem an der Ostküste vorangetrieben. Demgegenüber sei die Erprobung der Einsatzmöglichkeiten von Datenbankmanagementsystemen für die *Humanities* zunächst eher ein Westküstenphänomen gewesen (vgl. Liu 2008: 210 und 218).

14 | Neben Tunncliffe erwähnt Goldfarb auch Stanley Rice und Norman Scharpf als wichtige Wegbereiter deskriptiver Markupsprachen (vgl. Goldfarb 1995: 567f.).

15 | In der Praxis ist diese Unterscheidung von Primärinformationen und Metainformationen nicht trennscharf. So können den Elementen Attribute beigefügt werden, deren Werte Primärinformationen sein können. Demgemäß ist es möglich, die von einem Element umschlossenen Informationen als Attributwert in dem Element zu deklarieren. Eine Überschrift `<headline> ... </headline>` kann daher auch wie folgt annotiert werden `<headline text=„...“/>`.

(z.B. Überschriften, Absätze etc.) in konkreten Anwendungskontexten dargestellt werden sollen, kann infolgedessen gesondert spezifiziert werden. Inhalt und Aussehen werden zu Variablen, die (weitgehend) unabhängig voneinander verändert werden können.¹⁶

Die beschriebene Trennung von Inhalt und Form ist spätestens seit der Einführung von *Cascading Style Sheets* (CSS) durch das *World Wide Web Consortium* (W3C) ein zentrales Gestaltungs- und Entwicklungsprinzip von Webseiten. Denn mit der Auslagerung von Darstellungsanweisungen aus HTML-Dateien in CSS-Dateien wird die Handhabung von Inhalten und Formen weitgehend voneinander entkoppelt.¹⁷ Dabei treten, wie Liu herausgestellt hat, unterschiedliche Ebenen der Autorschaft zunehmend auseinander: das *Content Management* (Inhalt) und das *Consumption Management* (Form) (vgl. Liu 2008: 214f.).¹⁸

In der heutigen digitalen Medienpraxis sind die strukturierte Speicherung von Informationen in Datenbanken sowie die deskriptive Strukturierung von Dokumenten durch Markup zwei komplementäre Strategien zur Verwaltung und Distribution digitaler Informationen. Aus Datenbankinformationen können XML-Dokumente erzeugt werden und derartige Dokumente lassen sich in spezifischen

16 | Jedes Dokument kann also auf vielfältige Weise dargestellt und jede Darstellungsweise auf beliebig viele Dokumente angewendet werden; siehe hierzu das 2003 gegründete Projekt *CSS Zen Garden*, www.csszengarden.com/ (zuletzt aufgerufen am 14.06.2013).

17 | Zwar erlaubt bereits HTML die Auszeichnung von Dokumentstrukturen, doch vor der Einführung von CSS mussten die Darstellungsanweisungen innerhalb derselben Datei ausgewiesen werden, wodurch die formale Trennung von Inhaltsform und Darstellungsform praktisch unterlaufen wurde. Zudem war die Ausdrucksmächtigkeit von HTML hinsichtlich der Gestaltungsmöglichkeiten von Webseiten relativ gering, sodass HTML-Elemente zur Beschreibung von Inhalten häufig für gestalterische Zwecke »missbraucht« wurden. Dementsprechend diente das `<table>`-Element in der Frühzeit des WWW häufig nicht der Auszeichnung tabellarischer Inhalte, sondern dem Layout einer Webseite. Dies änderte sich mit der Einführung von Stylesheets, welche nicht nur die Spezifizierung typographischer Merkmale in einem Dokument ermöglichen, sondern auch die Positionierung und Anordnung von Informationsblöcken. Doch auch beim Einsatz von Stylesheets wird eine strikte Trennung von Inhalt und Form häufig unterlaufen. Für Gestaltungszwecke wird dabei nicht länger das Tabellenelement »missbraucht«, sondern das Gliederungselement `<div>`. Gemäß der Spezifikation von HTML (Version 4.01) soll dieses Element der inhaltlichen Gliederung von Dokumenten dienen (vgl. W3C 1999: Abschnitt 7.5.4). In der Praxis werden `<div>`-Elemente jedoch häufig aus rein gestalterischen Gründen eingesetzt. Hieran wird deutlich, dass die Separierung von Inhalt und Form ein Ideal ist, welches wahrscheinlich nie vollständig erreicht wird.

18 | Vom *Content Management* und vom *Consumption Management* unterscheidet Liu zudem das *Transmission Management*; siehe hierzu Liu (2008: 214).

Datenbanken speichern und verwalten. Die Datenbank als latente Infrastruktur geht jedoch über die auf Stylesheets und Markup basierende Trennung von Inhalt und Form hinaus. Mit der Verbreitung von datenbankbasierten Content Management Systemen (CMS) seit Anfang der 2000er Jahre setzte sich das Template-Paradigma beim Entwurf von Webseiten durch. Dabei kann durch die Speicherung von Webseiteninformationen in Datenbanken (und nicht in HTML-Dateien) nicht nur das Aussehen von Dokumenten (typographische Gestaltung und Anordnung der Inhaltselemente) manipuliert werden, sondern auch, welche Inhalte erscheinen sollen.¹⁹ Die an einen Nutzer gesandte Webseite (HTML-Dokument) wird beim Aufruf der URL dynamisch vom CMS erzeugt. Hierin besteht – vereinfacht ausgedrückt – die entscheidende Neuerung von CMS. Der *Content* der Webseite wird in einer Datenbank gespeichert, aus der diese Inhalte gemäß einer Vorlage – dem Template – ausgelesen und in von Browsern interpretierbare, d.h. darstellbare, HTML-Dateien übersetzt werden. Das Template spezifiziert einerseits, wie bestimmte Informationen (Überschriften, Fließtext, Hyperlinks etc.) dargestellt werden sollen. Andererseits wird im Template festgelegt, welche Informationen (Tupel, Attribute etc.) aus der Datenbank ausgelesen und wie diese Informationen auf der Webseite arrangiert werden sollen.²⁰ Darüber hinaus wird durch CMS die Personalisierung von Webangeboten ermöglicht, indem externe Variablen, wie z.B. der Ort des Webseitenaufrufs oder das Nutzungsverhalten, in die Auswahl der Datenbankinhalte einbezogen werden.²¹ Die Datenbank erscheint dabei aus Sicht der Endnutzer nicht als selektierbarer Informationsbestand, sondern bleibt weitgehend verborgen. Sie zeigt sich nur mittelbar, wenn beispielsweise unterschiedliche Zugriffsformen (Nutzer, Geräte, geographische Regionen) auf dasselbe Webangebot verglichen werden.

Auf der funktionalen Ebene zeigt sich die Datenbank als latente Infrastruktur in der Flexibilität des Umgangs mit Informationen, die neue Ausdrucks- und Gestaltungsmöglichkeiten eröffnet. Ein Beispiel hierfür sind Formen datenbankbasierter Erzählungen. Hierzu zählt unter anderen das bereits erwähnte Filmprojekt *Soft Cinema* von Lev Manovich und Andreas Kratky (2005).²² Bei diesem medienkünstlerischen Experiment mit neuen Formen des Filmschaffens wird die zeitliche (Reihenfolge) und räumliche (Platzierung im Split-Screen) Anordnung

19 | Derartige Selektionsmöglichkeiten sind der Markupsprache XML nicht immanent. Der selektive Zugriff auf XML-Dokumente wird durch an XML angelehnte Abfragesprachen wie z.B. XPath oder XQuery ermöglicht.

20 | Hierzu zählt unter anderem auch die dynamische Erzeugung von Navigationsmenüs, die sich bei Änderungen in der Webseitenstruktur automatisch auf allen betroffenen Seiten aktualisieren. Dies vereinfacht die (Weiter-)Entwicklung und Pflege von umfangreichen Webangeboten enorm.

21 | Auf die möglichen Nachteile der Personalisierung hat Eli Pariser (2011) hingewiesen, S. 266ff.

22 | Siehe hierzu S. 143f.

von Filmsequenzen automatisch von einer Software geleistet. Als Grundlage dieser Filme dient eine Datenbank, in der die einzelnen Filmsequenzen, angereichert mit Metadaten, gespeichert sind. Durch die Anpassung der Software bzw. der Parameter, mit der die Software Sequenzen selektiert und arrangiert, kann aus der Datenbank eine Vielzahl von unterschiedlichen Filmen erzeugt werden. Diese Form der Autorschaft charakterisiert Liu als parametrisierend (vgl. Liu 2008: 216f.).²³ Ein weiteres Beispiel ist die softwaregestützte Generierung von Berichten auf Grundlage von strukturierten Datenbeständen, wie sie die Firma *Narrative Science* entwickelt hat.²⁴ Zum Einsatz kam diese Technologie zuerst im Bereich der Sportberichterstattung. Aus erhobenen Daten über den Spielverlauf werden mithilfe komplexer *Artificial Intelligence*-Anwendungen automatisch journalistische Berichte generiert, die stilistisch zwar wenig elegant sind, über den Verlauf des Spiels jedoch korrekt informieren, wie Mercedes Bunz anmerkt (vgl. 2012: 13). Mittlerweile wird diese Technologie eingesetzt, um Daten aus den Bereichen Finanzen, Marketing und Werbung in Texte zu transformieren: Dem Leser wird nicht die Datenbank präsentiert, sondern ein Text, in dem die von Programmierern als bedeutsam eingestuften Inhalte der Datenbank zusammenfasst sind.

Diese algorithmischen Verfahren der Übersetzung von Datenbankinformationen in Texte oder Filme mit ihrer parametrisierenden Form der Autorschaft können in Anlehnung an Bernhard Rieder, der den Begriff wiederum von Paul Virilio entlehnt, als Sehmaschinen begriffen werden: »They are *vision machines* that not only extend our perception into the masses of information that would normally be far beyond human scope, but that also *interpret* the environment they render visible« (Rieder 2005: 29).²⁵ Bei den genannten Beispielen bleibt jedoch nicht nur die Datenbank,

23 | Liu spricht in diesem Zusammenhang von *data pours*, die automatisch mit Datenbankinformationen gefüllt werden: »data pours [...] are places on a page – whether a Web page or a word-processing page connected live to an institutional database or XML repository – where an author in effect surrenders the act of writing to that of parameterization. In these topoi, the author designates a zone, where content of unknown quantity and quality – except as parameterized in such commands as ›twenty items at a time‹ or ›only items containing *sick rose*‹ – pours into the manifest work from databases or XML sources hidden in the deep background« (vgl. Liu 2008: 216f.).

24 | Siehe hierzu die Webseite der Firma Narrative Science, www.narrativescience.com/ (zuletzt aufgerufen am 10.06.2013), sowie die Webseite des Projekts Stats Monkey an der Northwestern University, aus dem die Firma hervorgegangen ist, infolab.northwestern.edu/projects/stats-monkey/ (zuletzt aufgerufen am 10.06.2013).

25 | Rieder wendet das Konzept von Virilio auf Suchmaschinen an (vgl. Rieder 2005). Jedoch können auch algorithmische Verfahren der Übersetzung zwischen der Tiefe der Datenbank und verschiedenen Benutzeroberflächen als *vision machines* in dem genannten Sinn verstanden werden.

sondern auch das Verfahren der Sichtbarmachung unsichtbar. Gerade hierin besteht eine Gefahr, da der Endnutzer auf eine vordefinierte Perspektive festgelegt wird, ohne diese ändern zu können. Die Möglichkeit alternativer Narrative und konkurrierender Interpretationen wird durch die algorithmische Übersetzung unterlaufen, sofern man als Nutzer keinen Einfluss auf die Parameter der algorithmischen Übersetzung erhält. Entscheidend ist dementsprechend weniger, dass Datenbankinhalte anhand bestimmter Parameter selektiert und arrangiert werden, sondern wer auf welcher Ebene über diese Parameter entscheidet. Verlagert man diese Entscheidung auf die den Endnutzern zugewandten Benutzeroberflächen, dann tritt die Datenbank *als* Datenbank in Erscheinung, d.h. als Horizont der Selektion, Kombination und Interpretation von Informationen.

Die Operationsweise der Datenbank als latente Infrastruktur wurde bislang auf zwei Ebenen diskutiert: dem Aufkommen des Template-Paradigmas im Webdesign und in Bezug auf algorithmische Verfahren der Übersetzung von Datenbankinhalten in Texte oder Filme. Auf einer dritten Ebene ermöglichen Datenbanktechnologien aggregierende Verfahren der Zusammenführung, Auswertung und Präsentation von Inhalten aus verschiedenen Quellen in eigenständigen Webangeboten. Ein Beispiel hierfür sind die sogenannten *Social Media Dashboards* (z.B. HootSuite, Sprout Social, Netvibes), die eine einheitliche Schnittstelle zu vielfältigen Web 2.0-Services wie z.B. sozialen Netzwerken (z.B. Facebook, Google+, LinkedIn), Kommunikationsplattformen (z.B. Twitter), Blogs und anderen Web-Services (Flickr, YouTube, Foursquare) bieten. Als Metapher einer Kontroll- und Steuerungseinheit erlaubt ein Social Media-Dashboard die Beobachtung und Auswertung von Nachrichtenflüssen sowie die Publikation von Nachrichten auf verteilten Plattformen. Ein solches Dashboard ist ein Metaservice, »that curates our digital lives and adds value above the level of a single site« (Battelle 2011).²⁶

Voraussetzung hierfür ist die autonome Mobilität von Informationen in digitalen Netzwerken, die durch Programmierschnittstellen (APIs) zu den Datenbanken der einzelnen Anbieter gewährleistet wird. Durch die Bereitstellung von APIs erlauben Plattformen wie Facebook, Twitter etc. externen Entwicklern den Zugriff auf eigene Dienste und Informationsbestände, um auf dieser Grundlage Drittanbieter-Anwendungen zu erstellen. Vermittels APIs können Statusupdates, Tweets, Kommentare, Bilder etc. als Informationen aus unterschiedlichen Datenbanken ausgelesen bzw. darin abgelegt werden. Infolgedessen wird es möglich, mit denselben Informationen in verschiedenen Kontexten auf unterschiedliche Weise umzugehen. Hierdurch mag für die Nutzer der Eindruck der unbedingten Mobilität, Verarbeitbarkeit und Rekontextualisierbarkeit digitaler Informationen im WWW entstehen, doch diese Möglichkeit ist stets prekär. Denn durch technische oder juristische

26 | Als Beispiel eines Metaservices nennt Battelle die mittlerweile eingestellte Webseite *Memolane*, welche ihren Nutzern ermöglichte, ihre Aktivitäten auf verschiedenen Social-Media-Plattformen zu aggregieren und in einem Zeitstrahl zu visualisieren (vgl. Battelle 2011).

Einschränkungen einer API können die Möglichkeiten des externen Zugriffs auf Datenbestände jederzeit limitiert oder unterbunden werden.²⁷

ORIENTIERUNG IM VIELEN I: DAS EINE FINDEN

Fungiert die Datenbank als latente Infrastruktur, die aus der unsichtbaren Tiefe des Computers heraus die Weisen des Umgangs mit digitalen Informationen an der Oberfläche bedingt, dann bleibt diese als Informationssammlung für die Nutzer zumeist verborgen. In der medialen Praxis können Datenbanken jedoch auch *als* Datenbanken in Erscheinung treten, d.h. als Sammlungen von Informationen, die unsichtbar für die Augen der Nutzer in der Tiefe digitaler Speicher ruhen. Es handelt sich um ein Potenzial, das aktualisiert werden muss, indem Informationen in der Datenbank selektiert und an der Oberfläche zur Erscheinung gebracht werden. Die gespeicherten Informationen erscheinen hierbei als ein Bestand, auf den durch die Auswahl von Informationseinheiten zugegriffen und der algorithmisch ausgewertet werden kann.

Im Folgenden soll zunächst das Finden des Einen im Vielen in den Mittelpunkt gerückt werden, welches sowohl eine mediale Praxis als auch eine medienpraktische Herausforderung darstellt. Datenbanken fungieren dabei als Metamedien für das *close reading* medialer Konstellationen oder allgemeiner für *close practices*²⁸ mit medialen Konstellationen, die im Kontext der Datenbank als Informationseinheiten adressiert und selektiert werden können, d.h. von Webseiten, Bildern, Musikstücken, Videos, Personen, Profilen, Telefonbucheinträgen, Rezepten, Büchern, Produkten, Medikamenten, wissenschaftlichen Aufsätzen, bibliographischen Einträgen, Lexikonartikeln etc.²⁹ Bei der Suche nach dem *Einen* ist demzufolge nicht der gesamte Datenbankinhalt von Interesse, sondern einzelne in der Datenbank versammelte Elemente, an die sich unterschiedliche mediale Praktiken anschließen können, wie z.B. die (wissenschaftliche) Lektüre, die Kommunikation in sozialen Netzwerken durch Nachrichten, Kommentare u.ä., aber auch Kontroll- und Überwachungspraktiken.

Die Übergänge vom Vielen zum Einen sind genauer zu beschreiben. Zwei Fragen sind hierbei leitend: Erstens geht es darum, wie die in der Datenbank

27 | Ein Beispiel hierfür sind die zum Teil heftig kritisierten Änderungen, die Twitter an seiner API 2012 vorgenommen hat (vgl. Sippey 2012; Arment 2012; Caracciolo 2012).

28 | Der vorgeschlagene Begriff der *close practices* umfasst nicht nur die Lektüre oder Interpretation von medialen Konstellationen, sondern alle Formen des Umgangs *mit* und der Handhabung *von* medialen Konstellationen, wie z.B. das Bearbeiten, Teilen, Annotieren, Kommentieren etc.

29 | Die Bezeichnung Metamedium verweist hier auf eine spezifische Gebrauchsweise der medialen Konfiguration *digitale Datenbank*.

ruhenden Informationspotenziale an der Oberfläche aktualisiert werden, d.h. auf welche Weise »Informationsgewinnung organisiert wird« (Lovink 2009: 54). Zweitens ist danach zu fragen, wie die Informationspotenziale der Datenbank in ihrer Potenzialität an der Oberfläche zur Erscheinung kommen. Der Fokus liegt auf der Inszenierung der Datenbank als einem virtuell unerschöpflichen Informationsreservoir. Bereits 1998 hat Stephan Porombka auf diesen Aspekt von Datenbankschnittstellen hingewiesen. Er unterscheidet zwei Extrempole der Inszenierung dessen, »was als das Ganze auf Festplatten gespeichert« (Porombka 1998: 324) vorliegt: die *Flugperspektive* einerseits und die *Froschperspektive* andererseits. In Visualisierungen, wie z.B. Tree Maps, Netzwerken o.ä., könne die Komplexität des Ganzen in einer Überblicksdarstellung gezeigt werden, welche analog zu Bildern in der Flug- oder Vogelperspektive einen Blick von oben auf die Gesamtheit der Datenbank eröffnet. Was sich von dort zeigt, lässt sich Porombka zufolge jedoch »nur noch als komplexe Figur wahrnehmen, die für ihre eigene Komplexität einsteht« (Porombka 1998: 325). Auf der anderen Seite des Inszenierungsspektrums liege die »Froschperspektive ins Ganze« (Porombka 1998: 325), welche Porombka am Beispiel eines Knotens im Hypertext beschreibt.³⁰ Auf der Bildschirmoberfläche werde nur ein *Informationsbruchstück* präsentiert, das durch Anklicken von Links »durch weitere verbundene Informationsbruchstücke ersetzt werden« (Porombka 1998: 325) könne. Die Gesamtheit der potenziell zugänglichen Informationen bleibt zwar unsichtbar, doch werde das Potenzial der Datenbank erfahrbar. In Anbetracht der Verborgenheit des Vielen hinter dem Einen tritt das Ganze, wie Porombka darlegt, als Imaginäres zum Vorschein. Symptomatisch sei, »daß die Nutzer solcher Datenbanken immer sehr viel größere Datenmengen hinter der aufgerufenen Information vermuten, als tatsächlich abgespeichert sind« (Porombka 1998: 325).

30 | Implizit orientiert sich Porombkas medientheoretische Betrachtung von Datenbanken am Modell des Hypertexts. Diese Engführung wird vom Autor weder erläutert noch begründet. Doch tatsächlich lässt sich die hypertextuelle Struktur des WWW als eine Datenbank im nicht-technischen Sinn verstehen. Es ist ein Reservoir von Dokumenten (und von den in diesen enthaltenen Informationen), auf die mittels URLs zugegriffen werden kann. In der Tiefe der Computer ist das WWW global verteilt. An der Oberfläche jedoch ist es lokal, da alle Nutzer überall auf sämtliche Dokumente zugreifen können, sofern keine artifiziellen technischen Barrieren (Passwortschutz, Zensur, GeoIP-Sperren) sie daran hindern. Zwischen der globalen Tiefe der Speicher und den lokalen Benutzeroberflächen vermitteln Kommunikationsprotokolle (TCP/IP, HTTP, FTP usw.), denen jeweils ein Adresssystem (IPv4, IPv6, URL usw.) eingeschrieben ist, welches das Auffinden des Orts von Dokumenten im weltweit verteilten Speicherraum der Internetserver ermöglicht. Die Adressordnung des WWW ist gegenüber den Inhalten, die durch Aufruf einer Webadresse gefunden werden können, indifferent. Infolgedessen stellt die Orientierung im Web ein grundlegendes Problem dar.

Query: Anfragen an Datenbanken

Zentral für die Suche nach dem Einen im Vielen der Datenbank sind heute weder Überblick verschaffende Informationsvisualisierungen, noch das Mäandern durch hypertextuelle Netzwerke. Vielmehr ist die Suchanfrage (Query) der dominante Modus des Suchens und Findens in Datenbanken oder mittels Datenbanken, wie z. B. in Bibliothekskatalogen, Websuchmaschinen etc. Charakteristisch für diese Form der Orientierung in einer Informationssammlung ist nicht die Navigation durch einen Informationsraum, sondern die Formulierung von Suchbedingungen in einem Eingabefeld. Diese werden nach bestimmten Regeln automatisch mit dem Informationsbestand der Datenbank abgeglichen und in Ergebnisse übersetzt, die wiederum den Anfragekriterien entsprechen (vgl. Abb. 20). Die Query ist heute ein alltäglicher und normaler Bestandteil des Umgangs mit digitalen Medien.³¹ Ihre zentrale Stellung in der zeitgenössischen Medienkultur ist eng mit der rasanten Expansion des WWW und der damit einhergehenden Entwicklung von Websuchmaschinen verknüpft, die ein wichtiges Hilfsmittel sind, um Informationen im Web zu finden. Den entscheidenden Einschnitt stellte in diesem Kontext die Markteinführung der Suchmaschine Google dar, deren technische Funktionslogik, wie Röhle herausstellt, rasch zum Paradigma für Websuchmaschinen im Allgemeinen geworden ist: »Googles Version der algorithmischen Volltextsuche dominiert zurzeit den Zugang zu Online-Informationen vollständig« (Röhle 2010: 22).³²

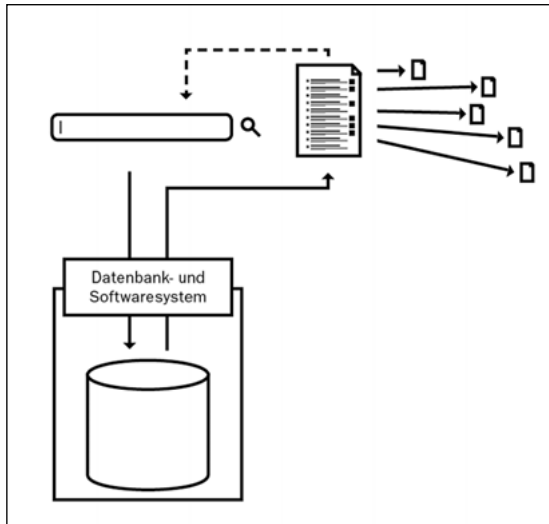
Die Sonderstellung der Query wird auf paradigmatische Weise in der minimalistischen Gestaltung der Startseite von Google deutlich, die bis heute nahezu

31 | Der Studie *Search Engine Use 2012* des Pew Internet & American Life Project zufolge verwenden 91% der US-Amerikanischen Internetnutzer prinzipiell Websuchmaschinen. Davon geben 59% an, Suchmaschinen täglich zu nutzen (vgl. Purcell et al. 2012: 5).

32 | Mit dem Erfolg der automatischen Indexierung von Webseiten durch die Google-Suche geht der zunehmende Bedeutungsverlust von Webverzeichnissen einher, die sich am Modell von Bibliothekskatalogen orientieren und eine redaktionelle Erfassung und Kategorisierung von Webseiten betreiben (vgl. Van Couvering 2008: 182f.; Röhle 2010: 22f.). Das kommerziell erfolgreichste Webverzeichnis der 1990er Jahre war *Yahoo!*, das 1994 von Jerry Yang und David Filo zunächst unter dem Namen *Jerry and David's Guide to the World Wide Web* gegründet wurde. Bereits 1998 begann *Yahoo!*, die Ergebnisse der Suche im Webverzeichnis mit Ergebnissen einer algorithmischen Suchmaschine zu ergänzen. Im Zuge dieser Entwicklung ist die Suchanfrage zum primären und privilegierten Modus der Suche im Web geworden, welche die Navigation durch die Kataloge der bis dahin dominanten Webverzeichnisse verdrängt hat. Zwar ermöglichen auch diese Suchanfragen im Katalog, doch aufgrund der Beschränktheit der Webverzeichnisse wurde bzw. wird bei der Eingabe von Suchworten allzu häufig nichts oder nichts Relevantes gefunden.

ausschließlich auf das Firmenlogo und das Suchformular reduziert ist.³³ Da die Webseite eigentlich keine Informationen enthält, werden die Nutzer dazu herausgefordert, Suchbegriffe in das Formularfeld einzugeben und so anfragend auf Googles Webdatenbank zuzugreifen. Zudem wird hierdurch suggeriert, dass etwas gefunden werden kann.

Abb. 20: Abstrakte Topologie einer Suchanfrage



Der von Websuchmaschinen vollzogene Übergang von der bibliothekarischen Logik der Katalogisierung zur algorithmischen Logik der Indexierung ist jedoch nicht zu generalisieren. Zwar haben sich Katalogisierungsverfahren als ungeeignet erwiesen, um der Größe, Vielfalt und Dynamik des gesamten WWW gerecht zu werden, in der digitalen Medienkultur kommt ihnen dennoch eine große Bedeutung zu. Ein Indiz hierfür ist die Vielzahl an konventionellen und zumeist relationalen Datenbanken, auf die durch das WWW zugegriffen werden kann (vgl. Madhavan et al. 2008: 1241f.). Diese beruhen, wie im Kapitel »Data + Access« dargelegt, nicht auf einer algorithmischen Zuschreibung von Bedeutung, sondern auf der Strukturierung von Informationen nach Vorgabe eines konzeptionellen Modells, welches Informationspartikeln eine Bedeutung zuweist. Der Unterschied bleibt an der Oberfläche jedoch weithin unsichtbar. Das Suchformular stellt sowohl bei Datenbanken im engen Sinn als auch bei Websuchmaschinen die primäre Form des Zugriffs dar. Aufgrund dieser Gleichförmigkeit der Suche in Informationssammlungen besteht an der Oberfläche zwischen Websuchmaschinen und Datenbanken im engen Sinn

33 | Ursprünglich war dieser Minimalismus keineswegs intendiert. Heute ist er jedoch ein zentrales Gestaltungsprinzip der Benutzerschnittstelle von Google (vgl. Röhle 2010: 154f.).

kein prinzipieller Unterschied.³⁴ Gesucht wird, indem Suchkriterien in ein Formular eingegeben werden. Auf Grundlage dieser Kriterien werden aus dem Informationsbestand der Datenbank automatisch Ergebnisse selektiert und an der Oberfläche zur Darstellung gebracht. Insofern lassen sich an der Benutzerschnittstelle allenfalls graduelle Unterschiede zwischen unterschiedlichen Anfragesystemen identifizieren. Diese betreffen die Art, wie Suchanfragen formuliert werden, die Möglichkeiten zur Spezifizierung von Anfragen sowie die Formen des Umgangs mit Suchergebnissen.

Abb. 21: Standardsuchformulare von Bing, Google und der Deutschen Nationalbibliothek

Quellen: www.bing.de, www.google.de, www.dnb.de

Häufig zeigen sich diese Differenzen jedoch erst in den erweiterten Suchfunktionen, die nicht Teil des *Default*-Interface sind, wie der Vergleich der Suchformulare von Google, Bing und der Deutschen Nationalbibliothek zeigt.³⁵ Der Unterschied zwischen der Funktionslogik von Websuchmaschinen und Bibliothekskatalogen verschwindet hinter der Homogenität der Suchinterfaces. Dies erweist sich jedoch mitunter als problematisch, z.B. wenn die verschiedenen Techno-Logiken

34 | In der Gleichförmigkeit der Suchinterfaces kann man im Anschluss an Winkler eine Naturalisierungsstrategie erkennen, welche die Erfüllung des Wunsches suggeriert, »tatsächlich über eine neutrale und transparente Erschließungsmaschine zu verfügen« (Winkler 1997b: 200).

35 | In diesem Zusammenhang hat Röhle darauf hingewiesen, dass es nicht allein entscheidend ist, »welche Einstellungsmöglichkeiten und Informationen dem Nutzer *theoretisch* zur Verfügung stehen, sondern [auch, M.B.] wie *zugänglich* diese sind« (Röhle 2010: 155).

unterschiedliche Interpretationen der gefundenen Ergebnisse erfordern. Um dies zu verdeutlichen, kann ein Beispiel herangezogen werden, mit dem Hendry und Efthimiadis (2008) aufzeigen, dass das Unwissen von Nutzern über die Funktionsweise von Websuchmaschinen unter Umständen falsche Interpretationen der Suchergebnisse zur Folge haben kann: Ein Nutzer wollte sich mithilfe von Google der Schreibweise des englischen Worts für »Kronleuchter« versichern und gab *chanda-leer* ein. Dass die Anfrage weniger als einhundert Ergebnisse hatte, interpretierte er fälschlicherweise als Bestätigung für die korrekte Schreibweise. Wie die Autoren berichten, hätte die Suche nach *chandelier* in der richtigen Schreibweise jedoch nahezu eine Million Ergebnisse zutage gefördert (vgl. Hendry/Efthimiadis 2008: 282).³⁶ Dass Google Ergebnisse zu einer Suchanfrage liefert, ist demzufolge nicht allein schon ein Indikator dafür, dass man den Suchbegriff richtig geschrieben hat. Hingegen wäre die Interpretation des Nutzers bei einer ähnlich gelagerten Suchanfrage an eine Katalogdatenbank nicht abwegig, da falsch buchstabierte Anfragen hier zumeist keine Ergebnisse haben.

Am Beispiel von Hendry und Efthimiadis wird deutlich, dass die richtige Interpretation von Suchergebnissen ein Wissen von der genutzten Suchtechnologie voraussetzt. Hierzu zählt unter anderem auch ein Verständnis davon, welche Ressourcen mit einem Anfragesystem durchsucht werden, d.h. von der Reichweite des Informationsbestands.³⁷ Daher ist es nicht nur problematisch, wenn Nutzer ein ungenügendes Wissen davon haben, wie Websuchmaschinen funktionieren, sondern auch, wenn sie meinen, dass die hinter den Suchformularen ablaufenden

36 | Mittlerweile verfügt Google über eine Rechtschreibprüfung, die in der Ergebnisliste einen Suchvorschlag zur Suche nach dem Wort in der korrekten Schreibweise anzeigt (»Meinten Sie: ...«). Infolgedessen ist es heute unwahrscheinlicher, dass ein Nutzer noch immer zu dem von Hendry und Efthimiadis beschriebenen Fehlschluss gelangt. Mit dieser technischen Erweiterung wurde das von den Autoren beschriebene Problem jedoch keineswegs gelöst, sondern verschoben: Wer heute nach etwas relativ Seltenem sucht oder einen Suchbegriff verwendet, der einer populären Anfrage sehr ähnlich ist, erhält von Google mitunter Vorschläge, die Suchanfrage zu verändern. Dies suggeriert dem Nutzer, einen Fehler gemacht zu haben, was jedoch nicht immer der Fall ist. Daher müssen die von Google unterbreiteten alternativen Suchvorschläge stets hinterfragt werden. Hierfür bedarf es eines Wissens über die Funktionsweise von Suchmaschinen, auf dessen Notwendigkeit Hendry und Efthimiadis mit ihrem Beispiel hinweisen.

37 | Mit der Einführung der sogenannten *Blended Search* oder *Universal Search* haben Websuchmaschinen spätestens seit 2007 damit begonnen, Ergebnisse aus verschiedenen Informationsbeständen (Nachrichten, Bilder, Videos etc.) in die normale Websuche zu integrieren (vgl. Quirnbach 2009). Hierdurch wird dem Nutzer einer Websuchmaschine tendenziell verborgen, welche Informationsbestände dieser durchsucht und wie die Ergebnisse aus den verschiedenen Beständen in eine Ergebnisdarstellung integriert werden.

technischen Prozesse immer denselben oder zumindest ähnlichen Prinzipien folgen.³⁸ Ein Indiz hierfür ist der generalisierende und undifferenzierte Verweis auf die Macht des Algorithmus, der einen beliebten Topos in aktuellen Debatten über die digitale Medienkultur darstellt (vgl. Passig 2012). Anstatt die Vielfalt der algorithmisch gesteuerten Informationsverarbeitungsprozesse im Computer differenziert zu betrachten, wird der Algorithmus als »undurchschaubare, orakelhafte« (Röhle 2010: 14) Macht mystifiziert.³⁹

Um die vermeintliche Allmacht von Algorithmen als Mythos zu entlarven und die tatsächliche Macht von Algorithmen besser zu verstehen, gilt es auch die Annahmen zu betrachten, die in die Gestaltung und Funktion von Anfrage-Interfaces einfließen. Zwischen den Extrempolen streng formalisierter, für Nutzer anforderungsreicher Anfragesprachen⁴⁰ und natürlichsprachlicher Anfrage-Interfaces⁴¹ dominiert in der heutigen Medienpraxis die Suche mit lose aneinander gereihten Schlagworten. Abgesehen von der Angabe der Suchbegriffe müssen weitere Suchkriterien dabei nicht mehr in einer formalen Anfragesprache expliziert werden. Diese sind vielmehr als implizite Vorannahmen in das Anfrage-Interface eingeschrieben. In der Standardsuche von Google werden Suchworte beispiels-

38 | Die Verdeckung der Unterschiede zwischen Suchtechnologien durch Suchinterfaces ist fraglos nicht die einzige Ursache für das genannte Problem. Zum Teil ist dieses auch auf das Desinteresse vieler Nutzer gegenüber der Funktionsweise der von ihnen verwendeten Technologien zurückzuführen.

39 | In den vergangenen Jahren wurde die Macht von Suchmaschinen in vielfältigen Beiträgen diskutiert. Wie Röhle in seiner Rekonstruktion der in dieser Debatte vertretenen Positionen herausarbeitet, wird in diesem Zusammenhang häufig das Bild eines (utopischen oder dystopischen) Determinismus gezeichnet, in dem die Macht der Algorithmen beschworen wird (vgl. Röhle 2010: 25ff.)

40 | Ein Beispiel hierfür sind die bibliographischen Information Retrieval-Systeme der 1970er und 1980er Jahre, die von ihren Nutzern zumeist die Kenntnis einer formalisierten Anfragesprache erforderten, welche es ihnen vor allem durch die explizite Verwendung boolescher Operatoren ermöglichte, komplexe Suchbedingungen zu formulieren, nach denen Informationen aus der Datenbank selektiert wurden (vgl. Borgman 1996: 493ff.).

41 | Das Ziel natürlichsprachlicher Anfrage-Interfaces ist, dass sich die Nutzer nicht den Anforderungen des Suchsystems anpassen müssen, sondern dass das System an die Sprache der Nutzer angepasst ist. Die technische Erfüllung dieses Versprechens wurde bereits 1996 von der Suchmaschine AskJeeves in Aussicht gestellt (vgl. Krajewski 2009; 2010: 149ff.). In der medialen Praxis konnten sich natürlichsprachliche Anfragesysteme bisher dennoch nicht durchsetzen. Heute finden sich natürlichsprachliche Anfragesysteme vor allem bei faktenbezogenen Informationen. Die *Wissensmaschine* Wolfram Alpha (www.wolframalpha.com) kann beispielsweise die Frage »What is the country with the fifth largest population?« richtig interpretieren und die korrekte Antwort »Brasilien« ausgeben.

weise per *Default-Einstellung* mit dem booleschen Operator AND verknüpft, d.h. es wird implizit davon ausgegangen, dass sich alle eingegebenen Suchworte in den Ergebnissen finden sollen. Zudem kommen insbesondere bei Websuchmaschinen algorithmische Verfahren der Interpretation von Suchanfragen zum Einsatz. Algorithmen, die auf Grundlage von erhobenen Nutzungsdaten entwickelt wurden, kategorisieren Suchfragen z.B. nach Genre, Thema, Intention, Ziel, Spezifität, Bandbreite, Autoritätsbezug sowie Orts- und Zeitbezug (vgl. Lewandowski 2011: 66f.).⁴² Bei dieser Form der Schlagwortsuche wird die explizite Angabe von Suchbedingungen zunehmend in die implizite Gestaltung des Anfrage-Interface und in die in der Tiefe unsichtbar ablaufenden Operationen des Suchsystems verlagert. Hierbei bleibt die Anfragelogik im bzw. hinter dem Interface verborgen. Wie Google und andere Websuchmaschinen Suchanfragen interpretieren ist für die Nutzer nicht unmittelbar ersichtlich, weshalb im Anschluss an Röhle bei der Analyse von Such-Interfaces auch danach zu fragen ist, ob und wie die impliziten Voreinstellungen und Interpretationsprozesse an der Benutzeroberfläche transparent gemacht werden (vgl. Röhle 2010: 156). Bisher lassen sich auf Seiten der Suchdienstanbieter jedoch kaum Bemühungen beobachten, die in diese Richtung weisen.⁴³

Zugleich hat die Verlagerung der expliziten Angabe von Suchkriterien in die implizite Gestaltung des Suchsystems zur Folge, dass die Einstiegshürde für Anwender relativ gering gehalten wird. Auch wenn derartige Anfragesysteme auf den inkompetenten Nutzer hin entworfen sind, verschwindet die Frage der nutzerseitigen Kompetenz nicht. Es bedarf zwar keines besonderen Könnens, um überhaupt etwas zu suchen, aber einer Kompetenz, um Suchanfragen zu stellen, mit denen gezielt bestimmte Informationen gefunden werden. Diese Kompetenz besteht nicht nur in der Wahl der Suchworte und erweiterter Suchoptionen, sondern auch in der Wahl einer geeigneten Datenbank und in der richtigen Interpretation der Suchergebnisse.

42 | Hierdurch wird, wie Lewandowski darlegt, die Zusammenstellung der sogenannten organischen Ergebnisse aus dem Webdatenbestand, die Integration von Ergebnissen aus den anderen Informationsbeständen (Bilder, Videos, Nachrichten, Bücher etc.), die Anfrageerweiterung, die Erzeugung von Suchvorschlägen sowie die Auswahl der anzuzeigenden Werbung beeinflusst (vgl. Lewandowski 2011: 59f.).

43 | Ein Indiz hierfür ist, dass Röhle in seiner Analyse des Such-Interfaces von Google keine Beispiele dafür anzugeben vermag. Wolfram Alpha (www.wolframalpha.com) stellt in diesem Zusammenhang eine erwähnenswerte Ausnahme dar, da die sogenannte *Wissensmaschine* die Interpretationen von Suchanfragen im Ergebnis-Interface sichtbar macht. Hierdurch wird der Kontext offengelegt, für den die als Ergebnis angezeigten Informationen zutreffend sind. Für mehrdeutige Anfragen werden ggf. sogar verschiedene Interpretationsmöglichkeiten aufgeführt, aus denen der Suchende auswählen kann.

Stream: Treiben im Informationsfluss

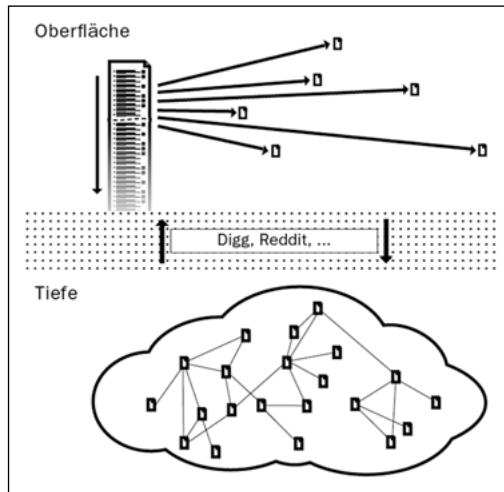
Von der Query als anfragender Zugriffsform auf Datenbanken unterscheiden Richardo Baeza-Yates und Berthier Riberio-Neto in *Modern Information Retrieval* das *Browsen* als zweite grundlegende Form der Interaktion mit Datenbankinformationen (vgl. 1999: 4). Der Nutzer wird beim Browsen nicht einem Such-Interface gegenüber verortet, sondern in einem Informationsraum, in dem er sich orientieren muss.⁴⁴ Paradigmatisch hierfür ist das Durchstöbern einer weitverzweigten Webseite nach bestimmten Informationen. Der Nutzer erscheint als ein navigierendes Subjekt, welches sich anhand von (hierarchischen) Menüstrukturen, *Breadcrumb*-Navigationen,⁴⁵ Sitemaps usw. orientieren kann. Mit dem Aufkommen des Web 2.0 gewinnt neben der Query und dem Browsen der Stream zunehmend an Bedeutung, um das Eine im Vielen zu finden. Der Stream, als eine an Aktualität und Popularität orientierte Präsentationsform des Vielen, ermöglicht keine gezielte themenspezifische Suche nach Informationen im WWW oder auf einer Webseite. Im Vordergrund steht vielmehr das relativ ungezielte Entdecken von neuen, aktuellen, populären und interessanten Ressourcen in bestimmten Themengebieten auf Grundlage von nutzergenerierten Inhalten, abgegebenen Bewertungen sowie anderen Nutzer- und Nutzungsinformationen. Eben dies machen sich Social-News-Aggregatoren (z.B. *Digg* und *Reddit*) und Social-Discovery-Tools (z.B. *StumbleUpon*) zur Aufgabe, welche auf der Logik des Streams beruhen. Aber auch der *Newsstream*

44 | Voraussetzung für das Browsen in einer Datenbank im engen Sinn von DBMS ist die Übersetzung des Informationsreservoirs in eine an der Oberfläche durchquerbare Form, wie z.B. eine Webseite. Die Inhalte der Datenbank werden hierbei entlang bestimmter Kategorien gruppiert und in einer hierarchischen oder polyhierarchischen Navigationsstruktur angeordnet, sodass die Nutzer die Datenbank an der Benutzeroberfläche navigierend erkunden, durchstöbern und filtern können. Die Übersetzung von Datenbankstrukturen in Webseiten findet sich beispielsweise bei Online-Händlern wie Amazon, die auf diese Weise ihre Produkte präsentieren. Gemeinhin werden zwei Verfahren der Übersetzung von Datenbanken in Navigationsstrukturen unterschieden. Top-Down-Informationsarchitekturen geben eine Struktur vor, in die Datenbankinformationen eingeordnet werden. Bei Bottom-Up-Architekturen wird den Datenbankinformationen kein Platz in einer vordefinierten Ordnung zugewiesen. Vielmehr wird ihre Ordnung automatisch aus den in der Datenbank gespeicherten Metainformationen abgeleitet. Morville und Rosenfeld weisen darauf hin, dass der Bottom-Up-Ansatz gut funktioniert, wenn aus der Datenbank relativ homogene Unterseiten generiert werden sollen, wie z.B. bei Produktkatalogen. In anderen Kontexten erweist sich der Top-Down-Ansatz als produktiver. Deshalb werden in der Praxis häufig beide Verfahren eingesetzt (vgl. Morville/Rosenfeld 2006: 64ff.).

45 | *Breadcrumb*-Navigationen zeigen Nutzern an, wo sie sich im Informationsraum der Webseite befinden, z.B.: Hauptmenü/Untermenü_1/Untermenü_2/...

von *Facebook* birgt das Potenzial, im permanenten Fluss von Statusupdates etwas Neues zu entdecken.

Abb. 22: Orientierung mit Digg, Reddit und ähnlichen Social-News-Aggregatoren



Social-News-Aggregatoren präsentieren Verweise auf Ressourcen als allgemeine oder thematische Listen, wobei aktuelle und populärere Beiträge weiter oben erscheinen als andere. Durch eine Positiv- oder Negativbewertung kann jeder Nutzer die Popularität von Ressourcen steigern oder verringern und so die Rangfolge der Verweise in der Liste beeinflussen.⁴⁶ Beliebte Inhalte besetzen einen vorderen Platz in der Liste und werden langsamer von aktuelleren Beiträgen verdrängt.⁴⁷ Hierdurch steigt die Wahrscheinlichkeit, dass populäre Beiträge zeitlich versetzt von einer größeren Zahl anderer Nutzer zur Kenntnis genommen werden.⁴⁸ Zudem wird

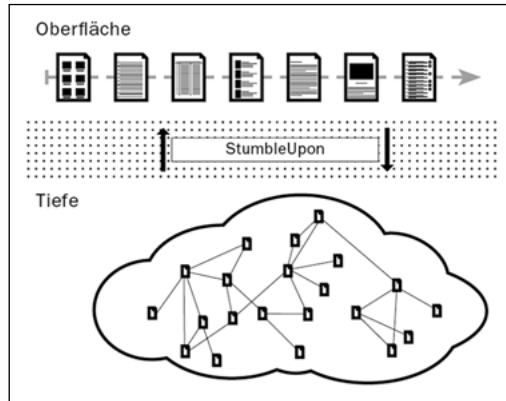
46 | Neben der binären Bewertung von Beiträgen können diese zumeist kommentiert und diskutiert werden. Diese Kommentare können wiederum durch die Nutzer bewertet werden. Zudem bietet Digg den Nutzern die Möglichkeit, anderen Nutzer zu folgen.

47 | Die Reihenfolge der Beiträge in der Liste wird algorithmisch berechnet, weshalb der Verdacht nahe liegt, dass neben Aktualität und Beliebtheit andere Faktoren in die Berechnung der Ordnung einfließen, die als redaktionelle Einflussnahme gewertet werden können. Insbesondere gegenüber Digg ist diese Vermutung häufig geäußert worden, weil der Ranking-Algorithmus nicht öffentlich gemacht wird. Der Quellcode von Reddit kann im Unterschied dazu im Internet frei eingesehen werden, <https://github.com/reddit/reddit> (zuletzt aufgerufen am 23.09.2013).

48 | Dies führt ähnlich wie bei Websuchmaschinen zu dem Problem, dass einzelne Nutzer versuchen, aus ökonomischen Interessen das Ranking gezielt zu manipulieren.

es möglich, Trends und ihre Entwicklung zu beobachten, um relevante Themen im WWW zu entdecken.⁴⁹

Abb. 23: Orientierung mit StumbleUpon



StumbleUpon beruht auf einem anderen Entdeckungsmechanismus. Nicht die Beobachtung von Trendlisten steht im Vordergrund, sondern das Umherstolpern im WWW ohne genaues Ziel, was vor allem durch Browsererweiterungen ermöglicht wird, die die *StumbleUpon* zur Verfügung stellt. Klickt man den *Stumble!*-Button des Plug-Ins, wird man automatisch zu einer Webseite weitergeleitet, die in der Datenbank von *StumbleUpon* verzeichnet ist. Diese kann der Nutzer lesen und bewerten. Aus diesen Bewertungen werden vor dem Hintergrund der Bewertungen anderer Nutzer automatisch Vorschläge errechnet, die man bei den nächsten Klicks des *Stumble!*-Buttons angezeigt bekommt.⁵⁰ Der Nutzer stolpert durch das WWW, wobei er – so das Kalkül – etwas finden kann, was für ihn von Interesse ist. Das zufällige, orientierungslose und unproduktive Umherirren im WWW wird in ein oberflächlich ebenso zufälliges, aber produktives Umherspringen übersetzt. *StumbleUpon* operiert im Hintergrund als unsichtbarer Empfehlungsdienst, der den Nutzer in einem linearen Stream von Vorschlägen situiert, dessen Ende nicht abzusehen ist. Hierdurch inszeniert sich *StumbleUpon* als ein virtuell grenzenloses Informationsreservoir. Daher ist die Erfahrung der Grenzen von *StumbleUpon* umso überraschender, die sich in der Nachricht manifestieren, dass zeitweilig keine weiteren Empfehlungen vorliegen.

49 | Die Möglichkeit, Trends identifizieren und beobachten zu können, ist nichts Spezifisches von Social News-Aggregatoren. Der Mikrobloggingdienst *Twitter* erzeugt beispielsweise aus den Tweets der Nutzer eine Liste von *Trending Topics*.

50 | Neben den von Nutzern erstellten Verweisen werden auch bezahlte Links in *StumbleUpon* aufgenommen, die als Werbung in den Empfehlungsstream eingewoben werden.

Streams folgen ähnlich wie die im Kapitel »Direct Access« (S. 206ff.) diskutierten Magnetbandspeicher einer linearen Zugriffs- und somit Präsentationslogik. Verteilt gespeicherte Nachrichten, Bilder, Videos, Blogeinträge usw. werden in einen linearen Informationsfluss gebracht, der das zufällige Entdecken des Einen im Vielen erlaubt. In dieser Hinsicht stellen Streams funktional keine Alternative zur Query als einem Modus der mehr oder weniger gezielten Suche nach Informationen dar. Wenn sich jedoch Nutzer mehr und mehr in den Nachrichtenflüssen von Facebook, Digg, Reddit, Tumblr usw. treiben lassen, drängt sich die Frage auf, wie Informationen in Streams organisiert werden und wer oder was entscheidet, welche Informationen wem gezeigt werden (vgl. Dash 2012). Diese Frage führt zurück zur technischen Logik der algorithmischen Ordnung von Informationssammlungen, wie sie im Kapitel »Data + Access« diskutiert wurde.

ORIENTIERUNG IM VIELEN II: DAS VIELE AUSWERTEN

In einer Sammlung von Vielem das Eine gezielt zu finden oder zufällig zu entdecken ist nur eine mögliche Erscheinungsweise der Datenbank *als* Informationssammlung. Von der im vorangegangenen Abschnitt beschriebenen Form der Orientierung im Vielen ist daher eine zweite Orientierungsform zu unterscheiden, die nicht darauf abzielt, eine im Bestand der Datenbank vorhandene, d.h. materiell verkörperte Information ausfindig zu machen. Der Informationsbestand der Datenbank kann auch als Basis für neue Informationen genommen werden. Die Datenbank dient in diesem Gebrauchskontext als Medium für das *distant reading* einer Informationssammlung (Moretti 2000: 57). Durch die kreative Auswertung von bekannten Informationen sollen noch unbekannte Einsichten geschöpft werden, welche es potenziell erlauben, Terroranschläge vorauszusehen und vorzubeugen, ökonomische Entwicklungen in Echtzeit nachzuvollziehen, den Ausbruch von Krankheitsepidemien zeitnah zu erkennen, Verbrechen zu bekämpfen, Freunde und Lebenspartner zu finden etc. Im Großen wie im Kleinen birgt die findige Auswertung vorhandener Informationen die Möglichkeit, neue Informationen zu erhalten. Insofern können digitale Datenbanken unter Umständen etwas wissen lassen, was so noch nicht gewusst, was allenfalls latent und rein virtuell als potentielle Information vorhanden war. Dies ist die Signatur des von Lyotard beschriebenen postmodernen Spiels vollständiger Informationen, welches sich in dieser Gebrauchsform von Informationssammlungen manifestiert.⁵¹

Bei der Auswertung des Vielen erscheint die Datenbank nicht als geschützte Aufbewahrungsstätte von Bekanntem, sondern als Grundlage oder Basis für das Entdecken von Unbekanntem. Diese Möglichkeit von bzw. Sichtweise auf Datenbanken findet ihren Ausdruck in der Metapher der Basis, welche im englischen Begriff *Data Base* enthalten ist. Die Überzeugung, dass aus der Auswertung von

51 | Siehe hierzu S. 183f.

Informationssammlungen neue Einsichten gewonnen werden können, hat die mediale Praxis mit digitalen Datenbanken von Anfang an begleitet.⁵² Unter dem Eindruck der leichten Verfügbarkeit immer größerer Informationsbestände kommt der Frage nach der Nutzbarmachung dieser Ressourcen durch analytische Auswertungsverfahren in der jüngeren Vergangenheit eine stetig wachsende Bedeutung zu. Bereits 1996 konstatierten Fayyad et al. einen Hype um das »knowledge discovery in databases« (Fayyad et al. 1996: 37), der sowohl in der Wissenschaft als auch in der Wirtschaft sowie der massenmedialen Berichterstattung zu beobachten sei. In der Aufmerksamkeit, die dieser Frage jüngst vor allem unter dem Stichwort Big Data zuteil wird, zeigt sich jedoch auch eine größer werdende Unzufriedenheit mit den etablierten Weisen des Umgangs mit Datenbanken sowie mit den tradierten Methoden der Auswertung von Informationsbeständen, welche Geert Lovink auf den Slogan gebracht hat: »Hört auf zu suchen. Fangt an zu fragen« (Lovink 2009: 63).

Im Folgenden soll zunächst der gegenwärtige Diskurs über die in Informationssammlungen ruhenden Wissenspotenziale beleuchtet werden. Daraufhin werden anhand von zwei Beispielen die Möglichkeiten und Grenzen der Auswertung des Vielen diskutiert. Der abschließende dritte Teil ist der Praxis und dem Diskurs der *Information Visualization* gewidmet, welche verspricht, unsichtbar in den Daten ruhende Zusammenhänge an der Oberfläche visuell erfahr- und erkundbar zu machen.

Auf der Suche nach dem Mehrwert

Nicht das Suchen nach dem Einen im Vielen ist Lovink zufolge in der heutigen Medienkultur entscheidend, sondern das Fragen danach, was aus dem Vielen gelernt werden kann. In eine ähnliche Richtung weist der programmatische Text »Against Search« von Manovich (2011), der die auf einzelne oder zumindest wenige Untersuchungsobjekte hin ausgerichteten Methoden der Geistes- und Kulturwissenschaften kritisiert.⁵³ Diese sind seines Erachtens nicht dazu geeignet, die Menge der digital verfügbaren *Kulturdaten* auszuwerten, weshalb neue Methoden

52 | Exemplarisch kann hierfür die von Horst Herold seit Anfang der 1970er Jahre vorangetriebene Computerisierung des Bundeskriminalamts angeführt werden, in deren Zuge die Rasterfahndung als eine Fahndungsmethode entwickelt wurde, die auf der Vernetzung und Auswertung von Informationsbeständen beruht (vgl. Gugerli 2006; 2009: 52ff.).

53 | Für Manovich stellt die Verfügbarkeit großer Informationsmengen eine Herausforderung dar, die seines Erachtens nicht mit den tradierten Methoden der Geistes- und Kulturwissenschaften gelöst werden kann: »The basic method of humanities and media studies which worked fine when the number of media objects were small – see all images or video, notice patterns, and interpret them – no longer works. For example, how do you study 167,00 images on Art Now Flickr gallery, 236,000 professional design portfolios on coroflot.com (both numbers as of 7/2011), or

und technische Verfahren für die Erforschung »der globalen digitalen Kulturen« (Manovich 2009: 221) nötig sind. Manovich nennt dieses Forschungsparadigma *Cultural Analytics* (vgl. Manovich 2009; Yamaoka et al. 2011). Im Kontext der *Digital Humanities*, deren Ursprünge bis in die 1940er Jahre zurückreichen, zirkulieren heute unterschiedliche Bezeichnungen für diese Art der Forschung. Von Michel et al. (2011) wurde die alternative Bezeichnung *Culturomics* vorgeschlagen. Die Autoren des 2011 in *Science* erschienenen Artikels definieren das hierdurch bezeichnete Forschungsfeld folgendermaßen: »Culturomics is the application of high-throughput data collection and analysis to the study of human culture« (Michel et al. 2011: 181). Im Mittelpunkt sollen quantitative Analysemethoden stehen, da diese nach Ansicht der Autoren zu präzisen Ergebnissen führen können.⁵⁴ Darüber hinaus wird die Notwendigkeit des Einsatzes computergestützter Auswertungsverfahren damit begründet, dass die verfügbaren Informationsmengen die menschlichen Kapazitäten bei Weitem überschreiten: »The corpus cannot be read by a human« (Michel et al. 2011: 176).

Jenseits der Kulturwissenschaften zeichnen sich auch in den Natur-, Lebens- und Sozialwissenschaften tiefgreifende Veränderungen hinsichtlich des Stellenwerts ab, den die automatische Auswertung großer Informationssammlungen einnimmt. In Reaktion hierauf proklamierte Chris Anderson (2008) in einem vielbeachteten, aber auch höchstumstrittenen Artikel das Ende der Theorie, welche von der Suche nach Korrelationen abgelöst werde:

»Petabytes allow us to say: ›Correlation is enough.‹ We can stop looking for models. We can analyze the data without hypotheses about what it might show. We can throw the numbers into the biggest computing clusters the world has ever seen and let statistical algorithms find patterns where science cannot.« (Anderson 2008)

Gegenüber dem im Magazin *Wired* publizierten Artikel wurde viel Kritik geäußert. Insbesondere Naturwissenschaftler wiesen in Reaktion auf Andersons Beitrag hin, dass das Entdecken von Korrelationen in Daten die Formulierung von Theorien und

176,000 Farm Security Administration/Office of War Information photographs taken between 1935 and 1944 digitized by Library of Congress« (Manovich 2011).

54 | Die Autoren des *Culturomics*-Beitrags haben ein szientistisches Wissenschaftsverständnis, auf dessen Grundlage sie die vermeintlich ungenauen qualitativen Methoden der Geistes- und Kulturwissenschaften diskreditieren. Ihres Erachtens konnten sich präzise quantitative Methoden in diesem Bereich bislang nur nicht etablieren, weil die verfügbare Datenbasis zu klein war: »Reading small collections of carefully chosen works enables scholars to make powerful inferences about trends in human thought. However, this approach rarely enables precise measurement of the underlying phenomena. Attempts to introduce quantitative methods into the study of culture have been hampered by the lack of suitable data« (Michel et al. 2011: 176).

Modellen nicht ersetzen könne, sondern diese allenfalls anzuleiten vermag. Es wird jedoch nicht in Zweifel gezogen, dass sich die Forschungspraxis mit der Verfügbarkeit und durch die Nutzbarmachung von stetig wachsenden Informationsbeständen verändert (vgl. Bollier 2010: 4ff.; Norvig 2009). Diese Verschiebungen exakt zu beschreiben erweist sich jedoch als Herausforderung.

Das Erkunden von Datenbanken mit dem Ziel, neue Informationen zu finden, verändert nicht nur die Wissenschaften, sondern auch die Wirtschaft und den Journalismus. Der Erfolg vieler führender Internetunternehmen, wie z.B. Facebook, Amazon, Google etc., basiert nach Ansicht von Mike Loukides zu einem Großteil auf der Sammlung und Auswertung von Information, wobei diese Unternehmen Daten auch dafür nutzen, Produkte aus Daten zu kreieren: »A data application acquires its value from the data itself, and creates more data as a result. It's not just an application with data; it's a data product« (Loukides 2010: 1). Ausgehend von dieser Beobachtung skizziert Loukides die Konturen einer *Data Science*, welche die Erstellung von ökonomisch verwertbaren Daten-Produkten unterstützen könne.

Der computergestützten Versammlung, Aufbereitung, Analyse und Veröffentlichung von Informationsbeständen kommt in den vergangenen Jahren auch im Bereich des (Online-)Journalismus eine wachsende Bedeutung zu. Für diese Entwicklung stehen Begriffe wie *Database Journalism*, *Data-driven Journalism* und *Datenjournalismus*.⁵⁵ Dabei ist der Einsatz von Computertechnologien bei der journalistischen Auswertung von Informationsbeständen keineswegs neu. Bereits seit Ende der 1960er Jahre kommen derartige Verfahren zum Einsatz (vgl. Meyer 2002: 192ff.). Als Bezeichnung war hierfür der Begriff des »computer-assisted reporting« (Meyer 2002: 79) gebräuchlich, ein Begriff, der in der Zeit der Mainframe-Computer geprägt wurde, mit Verbreitung des Personal Computers in den 1980er Jahren jedoch zunehmend an Schärfe und Aussagekraft verloren hat.⁵⁶

Ungeachtet der Frage, ob bzw. inwiefern der Datenjournalismus eine gänzlich neue Form der journalistischen Praxis darstellt, lassen sich spätestens seit 2009 Entwicklungen beobachten, die auf eine Konjunktur des Datenjournalismus hinweisen. Dies zeigt sich in den Bemühungen von national und international renommierten Tages- und Wochenzeitungen, datenjournalistische Formate in ihre Angebote zu integrieren.⁵⁷ Vorreiter hierfür waren vor allem die Tageszeitungen *The Guardian*

55 | Obwohl es viele Berührungspunkte und Überlappungen gibt, sind die Begriffe *Database Journalism* und *Data-driven Journalism* nicht gleichzusetzen. Der Fokus liegt hier vor allem auf dem *Data-driven Journalism*. Zum *Database Journalism* siehe exemplarisch Holovaty (2006).

56 | In der 2002 erschienenen vierten Auflage von *Precision Journalism* schreibt Meyer: »In a world where almost everything is computer assisted, that no longer means a lot« (Meyer 2002: 79). Die erste Auflage des Buchs ist bereits 1973 erschienen.

57 | Seit einigen Jahren unterhalten *The New York Times*, *The Guardian*, *Die Zeit* sowie andere Tages- und Wochenzeitungen speziellen Datenblogs, auf denen über

und *The New York Times*, welche sich unter anderem durch die interaktive Aufbereitung der von Wikileaks veröffentlichten Quellen über die Kriege in Afghanistan und im Irak hervorgetan haben (vgl. Matzat 2010). Ein wichtiger Bestandteil des Datenjournalismus ist neben der Berichterstattung auf Grundlage großer Informationssammlungen vor allem die Publikation von Originalquellen sowie die nutzerfreundliche Materialaufbereitung in interaktiven Visualisierungen. Hierdurch soll es den Nutzern dieser Angebote ermöglicht werden, Nachrichten nicht nur in Form von Artikeln oder Reportagen zu rezipieren, sondern sich selbst in den verfügbaren Ressourcen zu orientieren und diese zu erkunden (vgl. Rogers 2011).⁵⁸

Wissen und Unwissen der Datenbank

Die Suche nach den Informationen respektive dem Wissen, das bei der Auswertung von Informationssammlungen entdeckt werden kann, gewinnt rasant an Bedeutung. Motiviert ist dies gleichermaßen durch den Wunsch, sich in der Fülle der verfügbaren Informationen zu orientieren, sowie durch die Hoffnung, dass die Analyse von immer mehr Informationen zu einem besseren Verständnis unserer physischen, biologischen, gesellschaftlichen und kulturellen Welt führen wird. Ein erfolgreiches Beispiel hierfür ist die Früherkennung von Epidemien und Pandemien durch das World Wide Web, für die das von der kanadischen Gesundheitsbehörde in Kooperation mit der Weltgesundheitsorganisation entwickelte *Global Public Health Intelligence Network* (GPHIN) ein Vorreiter gewesen ist (vgl. Mawudeku/Blench 2005).⁵⁹ Das GPHIN verfolgt eine Strategie, Risiken für die öffentliche Gesundheit zu entdecken, die quer zu den traditionellen Mechanismen der Erkennung globaler Gesundheitsrisiken liegt.⁶⁰ Der hierarchisch organisierte

die jeweiligen Bemühungen im Bereich des Datenjournalismus berichtet wird. Ein weiterer Indikator für die wachsende Aufmerksamkeit, die dem Datenjournalismus zuteil wird, sind die zahlreichen Workshops und Konferenzen, die zu diesem Thema durchgeführt werden. Ein Überblick über diese findet sich auf datadrivenjournalism.net.

58 | Einen guten und empfehlenswerten Einblick zu den Entwicklungen im Bereich des Datenjournalismus gibt die Dokumentation *Journalism in the Age of Data* von Geoff McGhee aus dem Jahr 2010, datajournalism.stanford.edu/ (zuletzt aufgerufen am 12.03.2013).

59 | Nach zweijähriger Entwicklungszeit wurde die erste Version von GPHIN seit 1999 eingesetzt. Im Jahr 2004 wurde dieses System von GPHIN II abgelöst (vgl. Burns 2006: 769; Mykhalovskiy/Weir 2006: 43f.).

60 | Etwas Ähnliches leisten die 2008 eingeführten Google-Dienste *Flu Trends* (www.google.org/flutrends/) und *Dengue Trends* (www.google.org/denguetrends/). Die Vorhersagen beruhen hierbei auf der Auswertung der Häufigkeit von Suchanfragen zu einem bestimmten Thema (vgl. Ginsberg et al. 2009: 1012f.). Die Möglichkeit, Grippeepidemien auf der Grundlage von Suchanfragen zu identifizieren, hat zuvor

Informationsfluss von lokalen Gesundheitsbehörden zu regionalen, nationalen und internationalen Gesundheitsorganisationen wird vom GPHIN umgangen, indem es Seuchenfrüherkennung nicht durch die Untersuchung von Patienten, sondern durch die systematische Suche im World Wide Web betreibt (vgl. Mawudeku/Blench 2005: i9f.; Blench/Proulx 2005: 4f.).⁶¹ Seit 1999 beweist GPHIN, dass dem Mangel an offiziellen Gesundheitsinformationen der Überschuss an Informationen im Web produktiv gegenübergestellt werden kann. So warnte GPHIN bereits Ende 2002 vor einer ungewöhnlichen Atemwegskrankheit in China, mehrere Monate bevor die Weltgesundheitsorganisation eine offizielle SARS-Warnung herausgeben konnte (vgl. Blench/Proulx 2005: i8f.; Keller et al. 2009: 691). Bei GPHIN wusste man zwar nicht genau, um welche Krankheit es sich handelt, aber man wusste, dass etwas vor sich geht, was möglicherweise zu einem globalen Gesundheitsrisiko werden könnte.

Mit GPHIN wurde eine Möglichkeit gefunden, durch die Auswertung von im Web publizierten Berichten besser zu verstehen, was gegenwärtig in der Welt vor sich geht.⁶² Im größeren Kontext der digitalen Medienkultur hat diese Errungenschaft jedoch auch ihre Kehrseite, welche besonders deutlich zutage tritt, wenn aus Informationssammlungen nicht nur neues Wissen über die Vergangenheit und Gegenwart abgeleitet werden soll, sondern auch über die Zukunft. Exemplarisch kann dies am versuchten Bombenanschlag auf eine Maschine der Northwest Airlines am 25. Dezember 2009 verdeutlicht werden. An diesem Tag schmuggelte der damals 23-jährige Nigerianer Umar Farouk Abdulmutallab eine Bombe an Bord des Fluges NWA 253, die er kurz vor Erreichen des Zielflughafens Detroit detonieren lassen wollte (vgl. Delta Airlines 2009). Nur durch eine Fehlfunktion des Sprengsatzes und das Eingreifen von Crew und Passagieren scheiterte der Anschlag.

Obwohl letztendlich vereitelt, sind die Ereignisse dieses Weihnachtstages dennoch von besonderer Signifikanz, da im Anschluss an das versuchte Attentat eine sicherheitspolitische Debatte losbrach, die vor allem um die Frage kreiste, wie es möglich war, dass Abdulmutallab überhaupt an Bord des Flugzeugs gelangen konnte (vgl. Lipton/Shane 2009). Vor allem die Geheimdienste standen in der Kritik. Nicht weil sie keine Informationen über die Verbindungen von Abdulmutallab zu

bereits Eysenbach (2006) beschrieben, der derartige Analyseverfahren als *Infodemiology* bezeichnet.

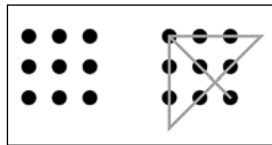
61 | Das GPHIN-System kombiniert automatische Suchroutinen mit menschlicher Expertise. Gefundene Berichte werden technisch akkumuliert und von Experten evaluiert, bevor Warnungen über Gesundheitsrisiken herausgegeben werden (vgl. Mawudeku/Blench 2005: i9f.). Die Überprüfung der Warnungen auf ihre Richtigkeit obliegt wiederum dem Global Outbreak Alert & Response Network (GOARN) der WHO (vgl. Mykhalovskiy/Weir 2006: 43).

62 | Auch für sogenannte Echtzeitanalysen bietet Google Dienste an, die hier exemplarisch genannt werden können: Google Trends (www.google.de/trends/) und Google Correlate (www.google.com/trends/correlate); siehe hierzu Choi/Varian (2009) sowie Mohebbi et al. (2011).

terroristischen Gruppierungen und einen möglichen Anschlag hatten, sondern gerade weil sie über diese Informationen verfügten. Das Problem bestand darin, dass man nicht in der Lage war, die verschiedenen Informationen zum Bild eines Terroristen und zur Prognose eines möglichen Anschlags zusammenzuführen. Zu diesem Schluss kommt der offizielle Untersuchungsbericht des Weißen Hauses. Das Hauptproblem lag demzufolge bei den Geheimdiensten, die daran scheiterten, einen Zusammenhang zwischen den Informationen herzustellen: »[T]he CT [Counterterrorism, M.B.] community failed to connect the dots« (White House 2010).

Anders als es der Bericht nahelegt, ist die Ursache für dieses Scheitern jedoch nicht allein auf die Nachlässigkeit der Geheimdienste bei der Auswertung und Analyse von Informationen zurückzuführen, sondern auch auf den Überschuss an Informationen, welcher die menschlichen Verarbeitungskapazitäten übersteigt und damit nur noch computergestützten Analysemethoden zugänglich ist. Diese Verfahren beruhen auf Algorithmen, die Informationsbestände gemäß bestimmter Regeln automatisch interpretieren. Hierauf gründen gleichermaßen die Chancen und Risiken der computergestützten Auswertung von Informationen. Algorithmen analysieren Informationssammlungen nach einem vorgegebenen Muster, wodurch diese in einen bestimmten Bedeutungskontext gestellt werden.

Abb. 24: Neun-Punkte-Problem



Um in Informationssammlungen Zusammenhänge zu erkennen und denselben Informationen einen neuen Sinn zu entlocken, ist es jedoch notwendig, mit einem bestimmten, vorgegebenen Kontext oder Rahmen zu brechen. Auf einfache Weise lässt sich dies an dem Neun-Punkte-Problem verdeutlichen, bei dem neun im Quadrat angeordnete Punkte durch vier Geraden miteinander verbunden werden sollen. Dieses Puzzle lässt sich nur dann lösen, wenn die Linien über die Grenzen des Quadrats hinaus gezogen werden, weshalb das Neun-Punkte-Problem als Sinnbild für die im englischen geläufige Redewendung *»thinking outside the box«* dient. Hieran scheitern Algorithmen, da sie nur im Rahmen ihrer eigenen Funktionslogik operieren können.

Auch wenn die Möglichkeit, dass die Zukunft mithilfe avancierter Computertechnologien vorhergesagt werden kann, einen beliebten Topos der Populärkultur darstellt, ist dies noch immer weitgehend Science Fiction. Daher wird man sich unter den Bedingungen der sicherheitspolitisch motivierten Zunahme an Überwachung künftig vermehrt mit dem Problem konfrontiert sehen, dass sich Ereignisse wie Terroranschläge schon immer in den Daten angekündigt haben werden. Das Wissen der Datenbank ist hierbei ein Wissen post Faktum im Konjunktiv II:

Man hätte es im Vorhinein wissen können. Damit ist eine Grenze der Auswertungsmöglichkeiten von Informationssammlungen benannt, die sich in Zukunft verschieben, aber wohl nie gänzlich verschwinden wird, da Computerprogramme stets nur »*inside their box*« operieren können.

Die diskutierten Möglichkeiten und Grenzen des Entdeckens von neuem Wissen in bereits bekannten Informationen führen zurück zu der bereits erwähnten Debatte über das beginnende Ende der Theorie unter den Bedingungen von Big Data. Dass Andersons These – die hypothesen- und theoriegeleitete Forschung werde durch die algorithmische Suche nach Korrelationen in großen Datenbeständen abgelöst – umstritten ist, wurde bereits erwähnt. Zugleich kommt der Korrelationsanalyse und anderen Data Mining-Verfahren gegenwärtig eine zunehmende Bedeutung zu. Daher argumentieren Viktor Mayer-Schönberger und Kenneth Cukier in ihrem Buch *Big Data: A Revolution That Will Transform How We Live, Work and Think*, dass das Fragen nach Kausalität im Vergleich zur Suche nach Korrelationen aktuell an Bedeutung verliert. Wichtiger als das *knowing why* (Kausalität) wird ihres Erachtens das *knowing what* (Korrelation): »Knowing *what*, not *why*, is good enough« (Mayer-Schönberger/Cukier 2013: 52).⁶³

Die Gegenüberstellung von Kausalität und Korrelation birgt die Gefahr, dass man beide als sich wechselseitig ausschließende Erkenntnismodi begreift. Dies ist jedoch nicht der Fall, denn Korrelationen nehmen in der natur- und sozialwissenschaftlichen Forschungspraxis schon lange eine wichtige Rolle ein. Kurzum: Korrelationen stehen nicht in Widerspruch zur Frage nach Kausalität und wissenschaftlicher Theoriebildung. Hierauf weisen auch Mayer-Schönberger und Cukier hin (vgl. 2013: 67f.). Was ist dann jedoch das spezifisch Neue an der Wissensproduktion im Zeitalter von Big Data? Für Mayer-Schönberger und Cukier besteht dies in der Rolle, die Analyseprogramme oder Algorithmen bei Suche nach Korrelationen einnehmen, welche in großen Informationsbeständen automatisch Zusammenhänge zwischen Variablen entdecken können (vgl. Mayer-Schönberger/Cukier 2013: 54ff.). Infolgedessen gehe das Finden von Korrelationen der wissenschaftlichen Hypothesen- und Theoriebildung voraus (Mayer-Schönberger/Cukier 2013: 55). Diese Behauptung ist insofern problematisch, als dass nahegelegt wird, dass Big Data-Auswertungsverfahren unbedingt und voraussetzungslos – oder in anderen Worten: theorie- und hypothesenfrei Zusammenhänge in Datensätzen entdecken können.

63 | Es ist bemerkenswert, dass Mayer-Schönberger und Cukier fast ausschließlich Beispiele aus ökonomischen Kontexten erwähnen, um ihre Diagnose zu untermauern: »Knowing why might be pleasant, but it's unimportant for stimulating sales. Knowing what, however, drives clicks. This insight has the power to reshape many industries, not just e-commerce« (Mayer-Schönberger/Cukier 2013: 52). Daher ist zu fragen, inwieweit der Korrelationsanalyse auch in den Wissenschaften eine größere Bedeutung zukommt als dem Fragen nach Kausalität.

Wie am Beispiel des Neun-Punkte-Problems gezeigt wurde, sind dem automatischen Entdecken von Zusammenhängen mit spezifischen Auswertungsverfahren jedoch stets Grenzen gesetzt. Daher gilt es zu fragen, welche Auswertungsverfahren aus welchen Gründen auf welche Datenbestände angewandt werden und wie die entdeckten Korrelationen interpretiert, verifiziert und nicht zuletzt theoretisiert werden. So mögen Algorithmen überraschende, d.h. nicht durch Hypothesen vorhergesagte Zusammenhänge zwischen spezifischen Variablen freilegen können, dass aber mit bestimmten Verfahren in bestimmten Datensätzen Korrelationen entdeckt werden können, bedarf eines Vorwissens, welches sich in die Auswertung einschreibt und die Interpretation der Ergebnisse bedingt. Es handelt sich, wenn man so will, um prozedurale Hypothesen, die von thematischen Hypothesen zu unterscheiden sind.⁶⁴ Während thematische Hypothesen einen Zusammenhang zwischen spezifischen Variablen behaupten, beruhen prozedurale Hypothesen auf der Annahme, dass mit spezifischen Verfahren bestimmte Zusammenhänge entdeckt werden können. Insofern markiert Big Data nicht den Anbruch einer Ära theorie- und hypothesenfreier Forschung. Vorläufig könnte zumindest eine zunehmende Verdrängung thematischer Hypothesen durch prozedurale Hypothesen diagnostiziert werden, welche der Analyse vorausgehen und sie anleiten. Diese prozeduralen Hypothesen sind den Auswertungsverfahren bislang weitgehend implizit. Aufgabe künftiger medientheoretischer Analysen von Big Data ist es daher, diese impliziten Annahmen zu explizieren.

Visuelles Erkunden mit Informationsvisualisierungen

Ziel der Auswertung von Datenbanken ist es Manovich zufolge, ein Verständnis von Informationssammlungen zu entwickeln sowie Muster und Zusammenhänge darin zu entdecken, »[to] understand the ›shape‹ of [the] overall collection and notice interesting patter[n]s« (Manovich 2011). Von zentraler Bedeutung sind die im Data Mining angewandten mathematischen Analyseverfahren. Fayyad et al. nennen sechs Methoden, die typischerweise zum Einsatz kommen: Klassifikation, Clusterbildung, Regression, Assoziationsanalyse, Anomalieerkennung und die Zusammenfassung respektive Verdichtung (vgl. Fayyad et al. 1996: 44f.). Obwohl es sich hierbei um automatische Auswertungsverfahren handelt, führen diese nicht

64 | Dieser Unterscheidungsvorschlag lehnt sich lose an Herbert A. Simons Unterscheidung von substanzieller und prozeduraler Rationalität an, die dieser im Kontext seiner Diskussion ökonomischer Theorien einführt. Verhalten beruht Simon zufolge auf substanzieller Rationalität, »when it is appropriate to the achievement of given goals within the limits imposed by given conditions and constraints« (Simon 1976: 131). Bei prozeduraler Rationalität stehen demgegenüber die Problemlösungsverfahren im Vordergrund: »From a procedural standpoint, our interest would lie not in the problem solution – the prescribed diet itself – but in the method used to discover it« (Simon 1976: 132).

zwingend zu neuen und gültigen Erkenntnissen. Die Ergebnisse des Data Mining bedürfen vielmehr der Beurteilung eines Interpreten. In diesem Sinn stellt das Data Mining den Versuch dar, in der für Menschen unverständlichen Gesamtmenge von Informationen Muster zu finden, die dem menschlichen Interpretationsvermögen zugänglich sind (vgl. Fayyad et al. 1996: 44). Einem ähnlichen Zweck dient die computergestützte Information Visualization. Dieser Ende der 1980er Jahre von Forschern am *Xerox Palo Alto Research Center* (Xerox PARC) geprägte Begriff bezeichnet einen Forschungszweig sowie eine mediale Praxis, die sich komplementär zu den mathematischen Analyseverfahren des Data Mining verhält (vgl. Mazza 2009: 8). Von der Visualisierung von Informationssammlungen mithilfe leistungsfähiger Computersysteme wird erwartet, dass sie dazu beitragen kann, das Problem der *Data Deluge* zu lösen (Fry 2008: 1).⁶⁵ Information Visualization erscheint als geeignetes Mittel, um nicht in der Datensintflut unterzugehen, sondern sich die verfügbaren Daten effektiv zunutze zu machen, indem man sie für Verstehen und Verständnis aufschließt: »Perhaps one of the best tools for identifying meaningful correlations and exploring them as a way to develop new models and theories, is computer-aided visualization of data« (Bollier 2010: 9).

Im Prozess der Visualisierung wird die Unanschaulichkeit, Unübersichtlichkeit und Unverständlichkeit von großen Informationssammlungen in die Anschaulichkeit, Übersichtlichkeit und Verständlichkeit von statischen oder interaktiven Schaubildern, Infografiken, Karten etc., d.h. von Diagrammen übersetzt. Hierdurch sollen die Informations- und Erkenntnispotenziale aktualisiert werden, die in den digitalen Informationsbeständen ruhen. Dieses Vermögen von Informationsvisualisierungen begründet Ben Fry wie folgt:

»The visualization of information gains its importance for its ability to help us ›see‹ things not previously understood in abstract data. It is both a perceptual issue, that the human brain is so wired for understanding visual stimuli but extends to the notion that our limited mental capacity is aided by methods for ›externalizing‹ cognition.« (Fry 2004: 33)

Indem Informationen in visuelle Reize übersetzt werden, wird das Denken Fry zufolge externalisiert. Ein ähnliches Argument bringt auch Ben Shneiderman vor, der das Leistungsvermögen der Informationsvisualisierung wie folgt fasst: »The process of information visualization is to take data available to many people and to enable users to gain insights that lead to significant discoveries« (Shneiderman 2006: VIII).

Neben einer Erkenntnisfunktion erfüllen Informationsvisualisierungen jedoch auch eine kommunikative Funktion. Hierauf weist der Schwede Hans Rosling, Pro-

65 | Ein Überblick über aktuelle Entwicklungen im Bereich der Informationsvisualisierung sowie vielfältige Beispiele für Visualisierungen finden sich auf den Webseiten *Information is Beautiful* (www.informationisbeautiful.net/), *Information Aesthetics* (infosthetics.com/) und *Visual Complexity* (www.visualcomplexity.com/vc/).

fessor für Internationale Gesundheit und eine der schillerndsten Figuren im Feld der Information Visualization, hin: »If the story in the numbers is told by a beautiful and clever image, then everybody understands.«⁶⁶ Visualisierungen sind demzufolge nicht nur Mittel des Forschens oder Erkennens, sondern auch Ausdrucksmittel von Erkenntnissen, die sich durch ihre leichte Verstehbarkeit und intuitive Nachvollziehbarkeit auszeichnen. Die Erkenntnisfunktion und die Kommunikationsfunktion von Informationsvisualisierungen stehen jedoch in einem problematischen Spannungsverhältnis zueinander, was im Folgenden dargelegt werden soll.

Einen wertvollen Ausgangspunkt eröffnet hierfür die Zeichentheorie von Charles Sanders Peirce. Bereits im Übergang vom 19. zum 20. Jahrhundert hat Peirce das Erkenntnispotenzial von Visualisierungen erkannt und im Rahmen seiner Zeichentheorie erörtert. Jedoch gebraucht Peirce (noch) nicht den Begriff der Visualisierung; er spricht vielmehr von Diagrammen, die er neben Bildern und Metaphern der Klasse der Bildzeichen zuordnet.⁶⁷ Da die Begriffe des Diagramms und der Informationsvisualisierung in aktuellen Debatten mitunter sehr heterogen gebraucht werden, sollten diese nicht ohne Weiteres gleichgesetzt werden (vgl. Wentz 2013: 202f.). Für die Auseinandersetzung mit Informationsvisualisierungen wird es sich jedoch als fruchtbar erweisen, diese als einen Typus diagrammatischer Darstellungsformen im Sinne von Peirce zu verstehen.

Die Besonderheit der Bildzeichen – d.h. von Bildern, Diagrammen und Metaphern – im Unterschied zu symbolischen und indexikalischen Zeichen ist nach Ansicht von Peirce, dass man aus deren Betrachtung – oder zumindest aus der Auseinandersetzung mit diesen – mehr Informationen über das Zeichenobjekt herausfinden kann, als in die Konstruktion des Zeichens hineingeflossen sind: »For a great distinguishing property of the icon is that by the direct observation of it other truths concerning its object can be discovered than those which suffice to determine its construction« (Peirce 1960: 2.279).⁶⁸ Auf paradigmatische Weise tritt dies an Diagrammen zutage, die im Unterschied zu Bildern und Metaphern nicht Qualitäten von Objekten zur Darstellung bringen (Bilder) oder Ähnlichkeiten zwischen Objekten herstellen (Metaphern), sondern Relationen zeigen. Diagramme sind Zeichen, die Strukturen, Funktionslogiken und Ereignisfolgen anschaulich machen. Ihre Besonderheit besteht darin, im Prozess der Veranschaulichung Möglichkeiten zur Rekonfiguration der diagrammatischen Zeichenkonstellation zu geben (vgl. Bauer/Ernst 2010: 46). Sie eröffnen ein Möglichkeitsfeld für Variationen,

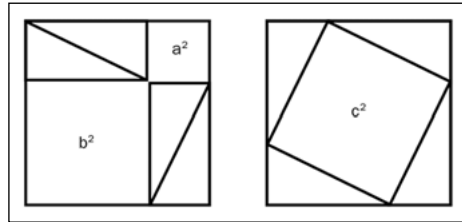
66 | Dies artikuliert Rosling in der 2010 produzierten Dokumentation *The Joy of Stats* (Hillmann 2010: 00:23:55).

67 | Die folgende Darstellung der peirceschen Diagrammatik orientiert sich an den Rekonstruktionen von Frederik Stjernfelt (2000) sowie von Matthias Bauer und Christoph Ernst (2010).

68 | Die *Collected Papers* von Peirce werden wie üblich in Dezimalnotation zitiert.

die im externalisierten Gedankenexperiment durchgespielt werden können.⁶⁹ Hierin liegt ihr Erkenntnispotenzial begründet.

Abb. 25: Beweis des Satz des Pythagoras



Quelle: Stjernfelt 2000: 369

Die aus Diagrammen ableitbaren Folgerungen sind Peirce zufolge deduktiv, d.h. es sind notwendige Schlüsse, deren Wahrheit sich im Diagramm zeigt.⁷⁰ Dies unterscheidet Diagramme grundlegend von den beiden anderen Arten von Bildzeichen: »It is, therefore, a very extraordinary feature of Diagrams that they *show*, – as literally *show* as a Percept shows the Perceptual Judgment to be true« (Peirce 1976: 318). Diagramme sind demzufolge Figurationen, an denen sich notwendiges Denken vollziehen kann, weil sie die Konsequenzen der Schlussfolgerungen evident vor Augen führen und dadurch deren Richtigkeit beweisen. Als Beispiel führt der Diagrammatikforscher Frederik Stjernfelt die Möglichkeit an, den Satz des Pythagoras zu beweisen, indem vier gleichförmige Dreiecke auf unterschiedliche Weisen in einem Quadrat angeordnet werden. Aus dem Vergleich der Anordnung der Dreiecke im linken und im rechten Quadrat wird deutlich, dass der Fläche von c^2 die Summe der Flächen von a^2 und b^2 entsprechen muss (vgl. Stjernfelt 2000: 369f.).

69 | Auch Stjernfelt spricht in diesem Zusammenhang von *Gedankenexperimenten*, in denen sich diagrammatisches Denken vollzieht (vgl. 2000: 369f.). In diesem Zusammenhang ist es wichtig darauf hinzuweisen, dass dieses Denken Peirce zufolge keine materialisierten Diagramme voraussetzt. Aus diesem Grund insistieren Bauer und Ernst darauf, das Forschungsfeld der Diagrammatik nicht auf Diagramme im engen Sinn von spezifischen verkörperten Bildzeichen zu begrenzen (vgl. 2010: 20f.). Für die Auseinandersetzung mit Informationsvisualisierungen ist hingegen die Möglichkeit der Externalisierung von Gedankenexperimenten von zentraler Bedeutung. Derartige Experimente sind aus dem Kontext der Wahlberichterstattung wohlbekannt. Die Sitzverteilung im Parlament wird zumeist als Kreisdiagramm oder Halbkreisdiagramm dargestellt. Anhand dieser Diagramme werden Parteienkonstellationen für mögliche Koalitionen durchgespielt.

70 | Dieser Aspekt ist von zentraler Bedeutung, wie Stjernfelt (2000: 369) und Bauer/Ernst (2010: 65f.) in ihren Rekonstruktionen der peirceschen Diagrammatik gleichermaßen herausstellen.

Diagramme sind Denkfiguren im buchstäblichen Sinn. An ihnen können in einem Prozess des sehenden Denkens notwendige Schlüsse vollzogen werden.⁷¹ Daher steht die Deduktion im Zentrum des diagrammatischen Folgerns. Für sich allein genommen sind die an Diagrammen vollzogenen Denkopoperationen rein formal. Sie geben keine Erkenntnisse über etwas anderes als über die Diagramme selbst preis. Visualisierungen beziehen sich jedoch auf Informationsbestände, die darüber hinaus zumeist auf ein Anderes Bezug nehmen, d.h. über etwas informieren. Infolgedessen steht nicht allein in Frage, welche Schlüsse aus Diagrammen gezogen werden können, sondern auch, ob durch ein Diagramm das Viele – in dem es sich zu orientieren gilt – auf adäquate Weise zur Darstellung kommt. Daher ist das diagrammatische Deduzieren nicht losgelöst von den abduktiven Hypothesenbildungen und induktiven Tests zu betrachten, in die das Denken an Diagrammen eingebettet ist (vgl. Bauer/Ernst 2010: 66f.). Diagrammatische Operationen vollziehen sich, wie Stjernfelt in seinen an Peirce anschließenden Überlegungen feststellt, in einem Regelkreis aus Abduktion, Deduktion und Induktion, weshalb dieses Denken im Erraten, Folgern und Überprüfen von Relationen und Zusammenhängen besteht (vgl. Stjernfelt 2000: 374).

Das Erstellen eines Diagramms gleicht einem Akt des abduktiven Ratens, der darin zum Ausdruck kommt, welche Zusammenhänge auf welche Weise veranschaulicht werden. Auf Grundlage dessen sind deduktive Schlüsse möglich, die wiederum induktiv am Referenzobjekt des Diagramms, d.h. an der Informationssammlung respektive einer wie auch immer gearteten Wirklichkeit überprüft werden müssen. Notwendig ist diese Überprüfung nicht zuletzt deshalb, weil sich die möglichen Folgerungen aus dem Typ des Diagramms ergeben, d.h. aus der Art und Weise, wie die Relationen, Strukturen oder Prozesse dargestellt werden, sowie aus der Syntax der erlaubten Transformationen des Diagramms und nicht aus den bezeichneten Zusammenhängen. Insofern ist es *gleichgültig*, aber nicht *gleichbedeutend*, ob man in einem Kreisdiagramm die Verteilung von Parlamentssitzen oder die Verteilung von Todesursachen veranschaulicht, da sich an dem Diagrammtyp Kreisdiagramm stets dieselben deduktiven Transformationen durchführen lassen.⁷² Der Typus konkreter Diagramme legt ihre Interpreten darauf fest,

71 | Mersch vertritt die Meinung, »dass alle diagrammatischen Visualisierungen, insbesondere aber Graphen, um interpretiert werden zu können, der Konventionalität und Regelhaftigkeit bedürfen. Keine Wissenschaftsvisualisierung kommt ohne Legende oder diskursiven Kommentar aus. Nicht nur verweisen Schrift und Bild im Diagrammatischen aufeinander, sondern die Diagrammatik selber erfordert den Text, der sie deutbar macht« (Mersch 2006c: 108). Im Anschluss hieran wäre zu fragen, ob die Regeln des deduktiven Folgerns mit Diagrammen ebenso auf Konventionen beruhen oder ob diese Diagrammen immanent sind. Eine Diskussion dieser weitreichenden Frage kann an dieser Stelle nicht geleistet werden.

72 | Sicherlich sind die Konsequenzen andere, die man aus diesen Schlüssen zieht, aber die deduktiven Folgerungen sind dieselben.

was sie an und mit diesen folgern können und damit auch, welche Erkenntnisse sie von diesen ableiten können.

Auf einen Begriffsvorschlag von Riemann rekurrierend, den sich auch Deleuze zueigen macht, können Diagramme als Mannigfaltigkeiten begriffen werden, welche eine dem »Vielen als solchem eigene Organisation« (Deleuze 1992a: 233) darstellen. Durch seine Bestimmung des Begriffs der Mannigfaltigkeit hat Riemann die Geometrie in der zweiten Hälfte des 19. Jahrhunderts vom Kopf auf die Füße gestellt. Hatte man bis dahin den Raum bei der Beschreibung geometrischer Objekte als transzendente Größe vorausgesetzt, macht sich Riemann dafür stark, Raum als eine geometrischen Objekten immanente Größe zu begreifen. Riemann zufolge können Objekte in Abgrenzung zur euklidischen Geometrie als Mannigfaltigkeiten verstanden werden, die nicht in Relation auf einen sie umgebenden und ihnen vorausgehenden Raum beschrieben werden. Geometrische Objekte werden vielmehr auf die ihnen innewohnende Räumlichkeit hin befragt. Jede Mannigfaltigkeit spannt demzufolge einen Raum auf und entfaltet dabei ihre eigenen Gesetzmäßigkeiten (vgl. Riemann 1990 [1854]).⁷³

Diese Eigenschaft wohnt auch diagrammatischen Darstellungsformen inne, die ihre Interpreten aus sich heraus auf eine Weltsicht festlegen, indem sie nur bestimmte Zusammenhänge, Relationen und Strukturen zur Darstellung bringen, an denen zudem nur spezifische Schlussfolgerungen vollzogen werden können. Hierauf hat Franco Moretti in Bezug auf die literaturwissenschaftliche Orientierung an Nationalliteraturen hingewiesen, welche seines Erachtens auf einem Denken in Baum-

73 | Riemann stellt kritisch fest, dass Raum als einer der Grundbegriffe der Geometrie noch nicht hinreichend geklärt sei. Für dieses Grundlagenproblem stellt das Konzept der Mannigfaltigkeiten eine Lösung dar. Als Konsequenz erweist sich der dreidimensionale euklidische Raum als kontingent, d.h. als möglicher Raum unter einer Vielzahl möglicher Räume: »Bekanntlich setzt die Geometrie sowohl den Begriff des Raumes, als die ersten Grundbegriffe für die Constructionen im Raume als etwas Gegebenes voraus. Sie giebt von ihnen nur Nominaldefinitionen während die wesentlichen Bestimmungen in Form von Axiomen auftreten. Das Verhältniss dieser Voraussetzungen bleibt dabei im Dunkeln; man sieht weder ein, ob und in wie weit ihre Verbindung nothwendig, noch apriori, ob sie möglich ist. Diese Dunkelheit wurde auch von Euklid bis auf Legendre, um den berühmtesten neueren Bearbeiter der Geometrie zu nennen, weder von den Mathematikern, noch von den Philosophen, welche sich damit beschäftigten, gehoben. Es hatte dies seinen Grund wohl darin dass der allgemeine Begriff mehrfach ausgedehnter Grössen unter welchem die Raumgrössen enthalten sind, ganz unbearbeitet blieb. Ich habe mir daher zunächst die Aufgabe gestellt, den Begriff einer mehrfach ausgedehnten Grösse aus allgemeinen Grössenbegriffen zu construiren. Es wird daraus hervorgehen, dass eine mehrfach ausgedehnte Grösse verschiedener Massverhältnisse fähig ist und der Raum also nur einen besonderen Fall einer dreifach ausgedehnten Grösse bildet« (Riemann 1990 [1854]: 304).

strukturen basiert. Von dieser Denkform ist die Beschäftigung mit Weltliteratur zu unterscheiden, die sich an der Denkfigur der Welle orientiert: »The tree describes the passage from unity to diversity: one tree with many branches [...]. The wave is the opposite: it observes uniformity engulfing an initial diversity« (Moretti 2000: 67).⁷⁴ Das Beispiel Morettis kann zu der Feststellung verallgemeinert werden, dass Diagramme, indem sie etwas zum Vorschein bringen, zugleich etwas verdecken. Sie verbergen die Kontingenz der ihnen immanenten Weltsicht. Oder anders formuliert: Diagramme machen stets nur bestimmte Zusammenhänge ersichtlich, aber gerade dies können sie nicht zeigen. Sie können in den Worten Dieter Mersch's »*die Modalitäten ihres Zeigens nicht mitzeigen* – sie verweigern die Sichtbarmachung ihrer Sichtbarmachung« (Mersch 2007: 64).

Diesem Problem wird im Kontext der computergestützten Informationsvisualisierung indirekt begegnet, insofern die Visualisierungspraxis nicht allein auf die Erstellung partikularer Diagramme abzielt, sondern auch auf die Entwicklung von generischen Visualisierungsanwendungen, welche es ihren Nutzern ermöglichen sollen, Informationen selbstständig in Diagrammen zu veranschaulichen und sie hierdurch zu erkunden. Information Visualization weist demzufolge über die in einer diagrammatischen Darstellung liegenden Möglichkeiten hinaus in Richtung des gesamten Prozesses der abduktiven Erstellung von Visualisierungen, der deduktiven Ableitung von Folgerungen an Visualisierungen und der induktiven Überprüfung dieser Schlüsse. Die Visualisierungssoftware wird als Schnittstelle zur Erkundung von Informationsbeständen begriffen, deren Möglichkeiten über die statischer Diagramme hinausgeht, da sie, wie Shneiderman anführt, die interaktive Einnahme unterschiedlicher Sichtweisen und Perspektiven auf Informationssammlungen ermöglicht:

»[A] great benefit of computing environments is the opportunity for users to rapidly revise the presentation to suit their tasks. Users can quickly change the rules governing proximity, linking, color, size, shape, texture, rotation, marking, blinking, color shifts and movements. In addition, zooming in or clicking on specific items to get greater detail increases the possibilities for designers and users [...]. A picture is often said to be worth a thousand words. Similarly, an interface is worth a thousand pictures.« (Shneiderman 2003: 373)

74 | Hintergrund der Gegenüberstellung von Bäumen und Wellen als zwei unterschiedlichen Denkformen bzw. -figuren ist Morettis eigenes Bestreben, Literatur nicht mehr nur auf der Ebene von Nationalliteraturen, sondern auf einem globalen Niveau zu analysieren. Er ist jedoch weit davon entfernt, das eine gegen das andere zu stellen: »Cultural history is made of trees and waves [...]. And as world culture oscillates between the two mechanisms, its products are inevitably composite ones« (Moretti 2000: 67).

Der von Shneiderman herangezogene Vergleich ist jedoch schief, denn Worte stehen zu Bildern nicht im selben Verhältnis wie Bilder – oder genauer Diagramme – zu Interfaces. Das Visualisierungsinterface stellt keine Alternative zu diagrammatischen Darstellungen dar, sondern eröffnet vielfältige Möglichkeiten zur Erstellung und Handhabung von zahllosen Visualisierungen.⁷⁵ Es ist ein Mittel zur Erkundung von Zusammenhängen, Mustern und Strukturen in Informationsressourcen. Dies lässt sich am Beispiel der *Gapminder World*-Software verdeutlichen, die auf Initiative von Hans Rosling von der Stiftung Gapminder entwickelt wurde (vgl. Gapminder).⁷⁶ Das Ziel der Software ist, verbreitete Mythen durch eine faktenbasierte Weltsicht zu entlarven: »Fighting the most devastating myths by building a fact-based world view that everyone understands« (Gapminder).⁷⁷

In *Gapminder World* sind die Daten von über 500 statistischen Indikatoren integriert, wie z.B. die Entwicklung der durchschnittlichen Lebenserwartung oder des Durchschnittseinkommens in einzelnen Ländern.⁷⁸ Diese statistischen Informationen können in animierter Form als Blasendiagramme oder auf einer Landkarte visualisiert werden. Die Möglichkeiten zur Erkundung der Informationsressourcen beschränken sich auf diese beiden Darstellungsformen. Variieren lassen sich ebenso die statistischen Indikatoren, deren Relation zueinander man über die Zeit hinweg betrachten kann, wie die Variable, auf die sich die Größe der Blase bezieht, sowie die Farbkodierung der Blasen.

Die kombinatorischen Möglichkeiten, die dies dem Nutzer eröffnet, sind enorm. Daher bleibt es dem Anwender der Software überlassen, sich diese zunutze zu machen und in den verfügbaren Informationen signifikante Zusammenhänge zu entdecken und Erklärungen für diese zu finden. Eben dies erfordert jedoch einen bestimmten Umgang mit den Visualisierungen sowie den zugrunde gelegten Daten. Um mithilfe von *Gapminder World* zu einer *faktenbasierten Weltsicht* zu gelangen, ist die alleinige Betrachtung einzelner Visualisierungen nicht hinreichend. Die am Interface erscheinenden Diagramme sind nur ein Bestandteil des Prozesses der Formulierung und Überprüfung von Hypothesen und Gegenhypothesen, bei dem bedeutsame Aspekte und Zusammenhänge von weniger bedeutsamen unterschieden

75 | Bemerkenswert ist zudem die Geschichte des Sprichworts »Ein Bild sagt mehr als tausend Worte«. Es handelt sich dabei keineswegs um ein altes chinesisches Sprichwort. Vielmehr entstammt es der Feder des Werbetexters Frederick Barnard, der das Sprichwort am 10. März 1927 in dem Branchenblatt *Printer's Ink* prägte, um das Werben mit Bildern zu bewerben (vgl. Safire 1996). Diese Geschichte kann als Indiz für den problematischen Status von Bildern als Erkenntnismitteln dienen.

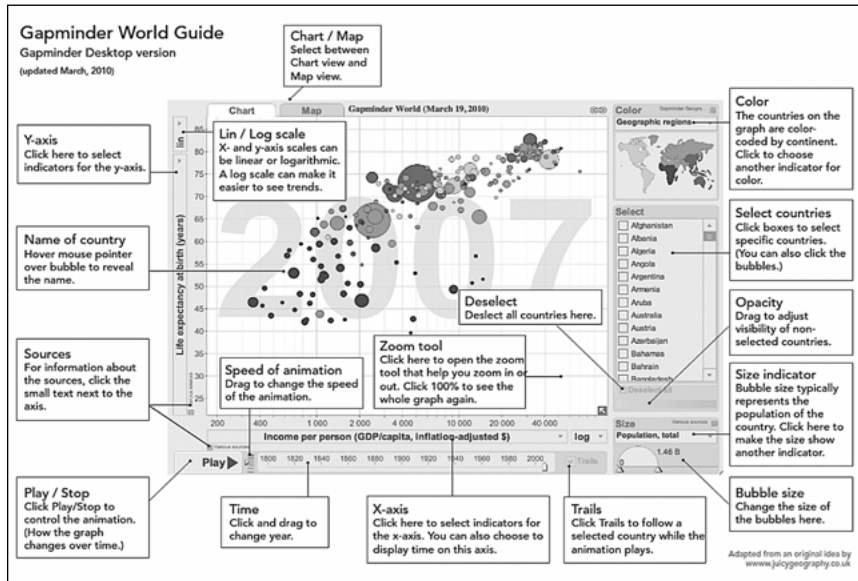
76 | Siehe hierzu die Webseite der Stiftung, www.gapminder.org (zuletzt aufgerufen am 10.02.2012).

77 | Weitere Beispiele für Visualisierungssoftware sind Tableau Public (www.tableausoftware.com/public) und Gephi (gephi.org/).

78 | Ein Überblick über die in *Gapminder World* integrierten Datensätze findet sich online, www.gapminder.org/data/ (zuletzt aufgerufen am 16.06.2013).

und unterschiedliche Perspektiven oder Interpretationen gegeneinander abgewogen werden müssen.⁷⁹ Daher erfordert die Visualisierungssoftware *Gapminder World* von ihren Nutzern eine (quasi-)wissenschaftliche Einstellung. Oder anders formuliert: Die Software verkörpert den Nutzer als forschendes Subjekt.

Abb. 26: Erläuterung des *Gapminder World*-Interfaces



Quelle: Gapminder 2010

Wie eingangs in Rekurs auf Rosling festgestellt wurde, erfüllen Informationsvisualisierungen nicht nur eine Erkenntnis-, sondern auch eine Kommunikationsfunktion. Vor allem dies stellt Rosling in seinen populären Vorträgen mit der Visualisierungssoftware zur Schau. Dabei steht die Frage, wie man mithilfe von *Gapminder World* zu neuartigen Erkenntnissen gelangen kann, im Hintergrund. Anhand der Visualisierungssoftware präsentiert Rosling vielmehr eine Erkenntnis, für die ihm eine Visualisierung als Beleg dient. So veranschaulicht er beispielsweise den Zusammenhang zwischen dem Wohlstand von Menschen – bemessen am durchschnittlichen Pro-Kopf-Einkommen in Ländern – und ihrer statistischen

79 | Wie David McCandless mit verschiedenen Visualisierungen des Verteidigungsbudgets unterschiedlicher Länder zeigt, ist eine der zentralen Fragen, ob absolute oder relative Größen zugrunde gelegt werden. So haben die USA im Jahr 2010 im weltweiten Vergleich zwar die größte Summe für Verteidigung aufgewendet. Relativ zum Bruttoinlandsprodukt war jedoch Myanmar Spitzenreiter (vgl. McCandless 2010).

Lebenserwartung (vgl. Rosling 2010).⁸⁰ Gezeigt wird hierbei, dass die durchschnittliche Lebenserwartung von Menschen mit dem gesellschaftlichen Wohlstandsniveau steigt. Die Visualisierung dieses Zusammenhangs in dem animierten Blasendiagramm von *Gapminder World* macht dies auf einfache und evidente Weise deutlich. Ob die Abhängigkeit der Lebenserwartung vom gesellschaftlichen Wohlstand hinsichtlich anderer Indikatoren zu relativieren ist, das kann man dem Diagramm allein jedoch nicht entnehmen.

Das Erkunden von Informationsressourcen mithilfe von Visualisierungen und die Kommunikation von Erkenntnissen durch Visualisierungen situiert diese in unterschiedlichen medialen Praktiken. Ist das Erkenntnispotenzial der Information Visualization in den Prozess der Visualisierung eingebettet – d.h. in den Prozess der abduktiven Erstellung, deduktiven Ableitung und induktiven Überprüfung von Visualisierungen, wie im Rekurs auf die peircesche Diagrammatik dargelegt wurde – so basiert das kommunikative Potenzial auf dem Resultat dieses Prozesses, nämlich der medialen Konstellation (Diagramm), dem eine bestimmte Gültigkeit zugesprochen wird. Als Beleg für die Richtigkeit des dargestellten Zusammenhangs ist ein Diagramm nicht hinreichend. Es bedarf, wie Mersch herausstellt, des sprachlichen Diskurses, der den Evidenzcharakter der Visualisierung begründet: »Das Bild beglaubigt sich allein durch den Text, und Evidenz existiert nur, wo Gründe gegeben sind, die sie als solche rechtfertigen« (Mersch 2005: 327).⁸¹

Aber dies wird beim kommunikativen Gebrauch von Informationsvisualisierungen tendenziell verdeckt, wenn der Prozess der Information Visualization nicht mitreflektiert wird. Dann scheint es, als könnten Visualisierungen selbst verbürgen, dass die im Diagramm erscheinenden Zusammenhänge ein Korrelat in der Wirklichkeit haben bzw. reale Sachverhalte adäquat abbilden. Hierin besteht der suggestive Selbstevidenzcharakter technischer Bilder, welche vermeintlich »das Ideal einer ›nichtintervenierenden‹ Objektivität« (Mersch 2005: 332) realisieren.

80 | Dieses Beispiel ist paradigmatisch, da es als Standardvisualisierung dient, die beim Öffnen der Software *Gapminder World* angezeigt wird.

81 | Bereits in den 1980er Jahren hat der Bildwissenschaftler Ernst H. Gombrich auf einen weiteren Aspekt der Beglaubigung wissenschaftlicher Bilder hingewiesen. Seines Erachtens ermöglichen es technische Bilder, »eine Unzahl von Fragen zu beantworten – allerdings immer unter der Voraussetzung, daß die technischen Spezifikationen des Instruments, Vergrößerung, Auflösevermögen usw., genau bekannt sind« (Gombrich 1984: 242). Gombrichs Fokus galt repräsentativen Bildern, aus deren Darstellung man Eigenschaften über das dargestellte Objekt rekonstruieren kann, sofern die Herstellungsbedingungen des Bildes bekannt sind. Gerade dies erweist sich bei Informationsvisualisierungen problematisch, da das dargestellte Objekt im Zuge der Visualisierung mitkonstituiert wird. Dennoch erweist sich das Wissen über die angewandten Visualisierungsverfahren sowie über die zugrunde gelegten Daten als notwendig, um Informationsvisualisierungen kritisch bewerten zu können.

Verstärkt wird diese Suggestion selbstevidenter Bezeugung sowohl ästhetisch als auch diskursiv. So weist Mersch kritisch darauf hin, dass Visualisierungen häufig an referenzielle Bildformen, wie z.B. Fotografien »angeähnelte« (Mersch 2006c: 110) werden. Auf der diskursiven Ebene wird der suggestive Selbstevidenzcharakter von Visualisierungen wiederum durch den argumentativen Stellenwert zementiert, der diesen zuerkannt wird.⁸² Hieran wird deutlich, dass das Erkunden von Informationssammlungen mit Visualisierungen und die Präsentation von Zusammenhängen in Visualisierungen in einem Spannungsverhältnis zueinander stehen. Für die Information Visualization erweist sich dies insofern als problematisch, als diese gleichermaßen verspricht, Erkenntnis- und Kommunikationsmittel zu sein.

Das Potenzial von Visualisierungen, Orientierung im Vielen digitaler Informationssammlungen herstellen zu können, soll hierdurch nicht grundsätzlich infrage gestellt werden. Vielmehr gilt es anzuerkennen, dass Informationsvisualisierungen weder neutrale noch selbstevidente Mittel sind. Die Anordnung von Informationen in diagrammatischen Strukturen entspringt nie aus der Datenbank selbst. Das Diagramm bzw. der Diagrammtypus gibt ein Ordnungsmuster vor, nach dem die Datenbankinformationen arrangiert werden. Dies ist stets zu bedenken. Insbesondere dann, wenn das Visuelle, wie Mersch herausstellt, nicht umhin kann, »als [...] die Evidenz einer Wirklichkeit zu suggerieren« (Mersch 2005: 341).

Hieran wird erneut deutlich, dass sich an die unterschiedlichen, in diesem Kapitel diskutierten medialen Praktiken mit Datenbanken verschiedene medien-theoretische Fragestellungen anschließen. Die Datenbank als latente Infrastruktur, das Auffinden des Einen in Informationssammlungen sowie die Auswertung des Vielen stellen unterschiedliche Gebrauchsformen von digitalen Datenbanken dar, die je eigene Phänomeno-Logiken entfalten, welche auf unterschiedliche Weise an die bereits diskutierten Techno-Logiken der Herstellung computer-lesbarer Signifikanz anschließen. Template, Dashboard, Query, Stream, Data Mining und Information Visualization lassen sich mithin nicht zu einer einheitlichen Logik digitaler Datenbanken zusammenführen. Sie wurden daher als heterogene mediale Praktiken mit Datenbanken beschrieben, in denen sich je eigene Logiken, d.h. Mikrologiken digitaler Datenbanken zeigen, an die verschiedene medientheoretische Problemkonstellationen anschließen.

82 | Kritisch auf diese Problematik hinweisend hat Julio Ottino 2003 in einem Artikel in *Nature* gefragt: »Is a picture worth 1,000 words?« (474). Ottino zeigt sich insbesondere gegenüber der von ihm beobachteten Zunahme von Bilder und Visualisierungen in wissenschaftlichen Publikationen skeptisch. Seines Erachtens stellen diese ein zu problematisierendes Ausdrucksmittel naturwissenschaftlicher Forschungen dar, dessen Gebrauch unter Regeln zu stellen sei, um dem naturwissenschaftlichen Streben nach wohlbegründetem Wissen weiterhin gerecht werden zu können.

