

4.2 NAHES UND FERNES WISSEN

Zum Glück jedoch folgt das WWW nur örtlich begrenzt der Logik der inventarisierten Bibliothek, deren Bände miteinander unverbunden sind. Weil ihre Einzelelemente miteinander verwoben sind, lassen sich in Hypertexten nicht nur die Textfragmente selbst als Ansammlungen von Codezeichen vermessen. Auch die zwischen ihnen bestehenden Verstrebungen lassen sich unter bestimmten Bedingungen auf ihre respektive ›Nähe‹ und ›Distanz‹ hin untersuchen. Stellen wir uns z.B. idealisiert ein virtuelles Museum vor, das eine Anzahl von digitalisierten Werken der abendländischen Kunst ausstellt. Die Sammlung sei um unseres Beispiels willen ein Querschnitt von der Renaissance bis ins 20. Jahrhundert. Organisiert sei dieses Museum als ein relativ konventioneller Hypertext, der jedem Ausstellungsstück genau eine Webseite widmet, welche Abbildungen, beschreibenden Text, vielleicht Videos und Tonmaterial zum Exponat enthält. Jede dieser Webseiten wäre also ein digitales Objekt im Sinne Yuk Huis.

Stellen wir uns nun weiter vor, diese einzelnen Webseiten wären untereinander verlinkt. Dies könnte auf verschiedene Arten geschehen, die es alle erforderlich machen, bestimmte diskrete Zeichen aus dem digitalen Gesamtobjekt hervorzuheben. So könnten z.B. ganz nach dem Wikipedia-Prinzip Worte im Beschreibungstext markiert werden, die direkt zu anderen Objekten oder Inventarseiten weiterleiten. Klickte man z.B. im Beschreibungstext eines Gemäldes auf den Namen eines anderen Gemäldes desselben Malers, würde man direkt zu diesem Werk weitergeleitet. Klickt man hingegen auf den Namen des Malers, würde man womöglich zu einer Übersichtsseite mit dessen Werken in chronologischer Reihenfolge gelangen. Ein Klick auf Begriffe wie ›Porträtkunst‹, ›Landschaftsmalerei‹, ›Kubismus‹ oder ›18. Jahrhundert‹ würde andere solche Überblicksseiten öffnen und somit eine assoziative Erschließung der Sammlung ganz im Sinne Vannevar Bushs ermöglichen – freilich mit dem Abstrich, dass die *trails of interest* hier nicht vom Nutzer erst geschlagen und ausgetrampelt werden müssten, sondern als virtuelles Straßennetz bereits vorhanden wären.

Auch auf Bildebene wäre eine solche Verlinkung durchaus denkbar – hier wäre es dann nötig, innerhalb der Abbildungen (den Ideen der Bense'schen Informationsästhetik folgend) diskrete Zeichen zu definieren und diese zueinander ins hypertextuelle Verhältnis zu setzen. So könnte z.B. die Auswahl eines Schädels oder einer Laute in einem Gemälde der altniederländischen Malerei zu Werken der selben Epoche weiterleiten, die das gleiche Symbol aufweisen, oder der Klick auf ein Christusmotiv epochenübergreifend Werke abrufen, welche die christliche Heilerzählung thematisieren. In diesem hypothetischen – von tatsächlich existierenden Angeboten aber durchaus nicht weit entfernten – virtuellen Museum ließe sich natürlich ganz leicht

topographisch vermessen, wie ›weit‹ einzelne Wissensgegenstände voneinander entfernt sind. Man könnte den Hypertext einfach als eine Landkarte oder ein Itinerar der Verlinkung darstellen.

Allerdings ließe sich dabei natürlich keiner Verbindungslinie eine klare Länge zuweisen, weil zwischen zwei Webseiten jeweils nur erhoben werden könnte, ob diese verbunden oder unverbunden sind, bzw. wie viele assoziative ›Schritte‹ unternommen werden müssen, um von Seite A zu Seite B zu gelangen. So ließen sich z.B. Aussagen treffen wie: ›Die kürzeste Verbindung zwischen Leonardo DaVinci und Jackson Pollock führt über sieben individuelle Links‹, oder: ›Es gibt keinen direkten Zusammenhang zwischen Norman Rockwells *The Problem We All Live With* und Caspar David Friedrichs *Wanderer über dem Nebelmeer*‹. Eine solche Topologie würde also die spezifische *connectedness* dieses bestimmten Hypertextes beschreiben und damit letztlich nur auf die kunstpädagogischen Prämissen und Intentionen schließen lassen, welche die Kuratoren bzw. Webdesigner zugrunde gelegt haben. Über die Interessenlagen der virtuellen ‚Besucher‘, oder gar die weiteren kulturellen Zusammenhänge, aus welchen diese hervorgehen, verrät eine solche Analyse wenig, weil sie die Perspektive der *Karte*, nicht die der *Strecke* einnimmt.

4.2.1 Soziometrie und Science Citation Index: Vorspiel zur Suchmaschine

Will man ein Netzwerk erfolgreich navigieren, so genügt es nicht, zu wissen, *wo* sich Querverbindungen befinden – wichtig ist darüber hinaus auch, zu wissen, *wie* diese beschaffen sind. Dies führt uns zurück zur bereits erwähnten Soziometrie Jakob L. Morenos. Morenos Versuch, soziale Beziehung kartier- und quantifizierbar zu machen, führte relativ schnell zu der Erkenntnis, dass die Netzwerkstruktur sozialer Gruppen zumeist keine gleichmäßige Verteilung von Personen und Beziehungen ist. Vielmehr finden sich in jedem Soziogramm Punkte (und damit Personen), um die herum sich die Beziehungsvektoren in besonderem Maße verdichten, die also ein besonderes soziales Kapital auf sich vereinen und die Moreno – in Anspielung sowohl auf ihre offensichtliche Popularität als auch auf das Aussehen der so vernetzten Punkte in der graphischen Darstellung – als *stars* bezeichnete (vgl. Mayer 2010: 68ff.). In den 1940er Jahren wandte sich die Soziometrie von der unübersichtlichen und potenziell verwirrenden Abbildungsform des Soziogramms zunehmend ab und ersetzte es durch Matrizen, welche den Vernetzungskoeffizienten einzelner Personen als diskreten numerischen Wert ausdrückten. Dabei entschied nicht nur die individuelle Vernetzung der betreffenden Person über ihren sozialen ›Wert‹, sondern auch die all jener Gruppenmitglieder, mit denen sie direkt oder indirekt verbunden war. Der *Status* eines Menschen in einem sozialen Gefüge ergibt sich demnach nicht nur aus

der reinen Anzahl seiner Bekanntschaften, sondern wiederum auch aus deren Gewichtung innerhalb der Gruppe (vgl. Mayer 2010: 71).

1964 schlug der amerikanische Linguist Eugene Garfield vor, diese Einsichten über die Natur von Netzwerken auf ein anderes Feld zu übertragen: Das der wissenschaftlichen Forschung bzw. genauer jenes der einander zitierenden wissenschaftlichen Veröffentlichungen. Garfield ging dabei von ganz ähnlichen Prämissen aus wie keine zwanzig Jahre zuvor Vannevar Bush in seiner Kritik an der Bibliothek: Angesichts unserer ständig und rasant schnell wachsenden Wissensbestände ließe sich die Forschungsliteratur unmöglich so schnell erschließen, wie sie gedruckt wird – unserem »scientific wealth« stehe, so Garfield, eine »bibliographic poverty« (Garfield 1984a: 525) gegenüber. Dies sei, so fährt Garfield fort, nicht zuletzt ein ganz schlichtes Problem der Forschungsökonomie: Jeder Index publizierter Forschungsarbeit sei in seiner Anlage mit einem erheblichen Arbeitsaufwand verbunden, und dieser sei nur zu rechtfertigen, wenn sein Nutzen seine Kosten überwiegt. Als eine naheliegende Gefahr schlechter Auffindbarkeit von Forschungsergebnissen nennt Garfield die Doppelforschung – also die unnötige Forschung an einem Gegenstand, zu dem bereits Untersuchungen vorliegen.

Allerdings seien in vielen Fällen die an redundante Forschung verschwendeten Ressourcen immer noch günstiger als die Kosten einer zentralisierten Erfassung aller Forschungsliteratur. Eine sinnvolle Indexierung wissenschaftlicher Veröffentlichungen müsse daher in erster Linie kostengünstig funktionieren – an eine tatsächlich beim Inhalt ansetzende Erschließung sei somit nicht zu denken (vgl. ebd.: 526). Stattdessen schlägt Garfield vor, sich den schon von Vannevar Bush diagnostizierten Netzwerkcharakter akademischer Literatur zunutze zu machen, um die Auffindbarkeit relevanter Schriftstücke zu gewährleisten. Sein bis heute in Gebrauch befindlicher (und in zahlreichen Forschungszweigen über Karrieren bestimmender) *Science Citation Index* nahm sich den existierenden juristischen Zitationsindex *Shepard's* zum Vorbild und sollte die Wichtigkeit von Büchern und Aufsätzen auf ähnliche Weise als mathematischen Wert ausdrücken, wie es in der Soziometrie mit dem sozialen Status von Menschen geschieht (vgl. Mayer 2010: 74f.).

Weil wissenschaftliche Publikationen sich laufend aufeinander beziehen und sich gegenseitig zitieren, lassen sie sich grundsätzlich genau so behandeln wie Menschen, die sich untereinander »kennen« – und genau wie man in einer Gruppe von Menschen besonders beliebte Mitglieder ausfindig machen kann, welche im Zentrum weitreichender Beziehungsgefüge stehen, findet man auch in Netzwerken von aufeinander Bezug nehmenden Schriftstücken solche, die für ihren Einzugsbereich zentral sind und hier das bilden, was wir umgangssprachlich »Standardwerke« nennen würden. Diese Werke, die innerhalb des Literaturnetzwerkes einen hohen Status aufweisen, bilden gewissermaßen Schlüsseltexte, denen zum einen ein hoher inhaltlicher Wert beigemessen wird und die dementsprechend oft zitiert werden, und die zum anderen

aus ihrer besonderen fachlichen Geltung heraus imstande sind, über ihre Literaturverzeichnisse und Fußnotenapparate als Wegweiser zu anderer wichtiger Literatur zu gelten (vgl. ebd.). Um diese Geltung als einen konkreten Zahlenwert ermittelbar zu machen, muss der Index nach Garfield aus zwei Listen bestehen: einer, die zitierende und einer, die zitierte Werke aufführt. Derselbe Text kann durchaus auf beiden Listen auftauchen – die Unterscheidung zwischen Zitieren und Zitiert-Werden ermöglicht es, Zitations-»Vektoren« zu bestimmen, entlang derer Veröffentlichungen einander Relevanz verleihen. So könnte ein Paper, das eine sehr große Zahl von Titeln zitiert, in einer topographischen Darstellung wie der des Soziogrammes durchaus wie ein *star* aussehen, tatsächlich aber völlig irrelevant sein, wenn es seinerseits in wichtigen Publikationen nicht als Referenz angegeben wird. Umgekehrt könnte ein Text, der für ein Forschungsgebiet Pionierqualitäten aufweist (wie dies z.B. für Vannevar Bushs *As We May Think* der Fall ist) einen enorm hohen Zitationsquotienten aufweisen, ohne selbst viele oder überhaupt irgendwelche Werke zu zitieren (vgl. Garfield 1984a: 528f.). Der *SCI* macht die Zentren und Peripherien wissenschaftlicher Diskurse anhand ihres Niederschlages in der Veröffentlichungspraxis ausfindig. Dabei sind nach Garfields Einschätzung ausdrücklich nicht nur die Zentren von Interesse. Ein Mangel an Zitationen deutet nämlich nicht immer auf ein schlechtes Paper hin – manchmal könne er auch ein Zeichen für Forschungsarbeit sein, die ihrer Zeit zu weit voraus und bei ihrem Erscheinen schlicht noch nicht anschlussfähig war. In Verbindung mit inhaltlichen Analysen könne die Wichtigkeitsvermessung, die der *SCI* vornimmt, dazu beitragen, unterbewerteten und zu Unrecht marginalisierten Arbeiten zur verdienten Aufmerksamkeit zu verhelfen (vgl. ebd.: 532).

Garfield selbst prägte den Begriff der *Hypersearch* als Bezeichnung für das Vorgehen, statt einer großflächigen Erschließung von Einzeltiteln jene zentralen Schlüsseltexte ausfindig zu machen, über welche man zu den relevantesten Titeln im betreffenden Themenfeld gelangt (vgl. Mayer 2010: 74f.). Und so revolutionär dieser Ansatz für die Wissenschaftsökonomie und -soziologie bereits gewesen sein mochte – sein eigentlicher, die gesamte Kultur der neuen Medien durchdringender Siegeszug sollte erst Ende der 1990er Jahre, mitten in der *take-off*-Phase des World Wide Web, beginnen.

4.2.2 Google: Kulturelle Relevanz als statistische Korrelation

Angeichts der Fülle der bereits existierenden und leicht verfügbaren Literatur zum Aufstieg des *Google*-Imperiums scheint es hier wenig sinnvoll, die Geschichte des Unternehmens im Detail zu rekapitulieren. Vielmehr soll im Folgenden am Beispiel *Googles* abgehandelt werden, wie das Suchparadigma der *Hypersearch* zu einer entscheidenden Kulturtechnik der virtualisierten Welt werden konnte – und wie sich diese Kulturtechnik auf die Inhalte auswirkt, die mit ihrer Hilfe organisiert, gefunden

und rezipiert werden. Im weiteren Verlauf der Arbeit wird dann der Blick auf die Frage zu richten sein, was genau dies für das mögliche Funktionieren und Nicht-Funktionieren musealer Angebote im Netz bedeutet.

Im Jahre 1998 veröffentlichten die an der Universität Stanford tätigen Computerwissenschaftler Sergey Brin und Larry Page ein Paper mit dem recht spröden Titel *The Anatomy of a Large-Scale Hypertextual Web Search Engine*. Die von Brin und Page darin in Aussicht gestellte Suchmaschine trug zu diesem Zeitpunkt bereits den Namen *Google* und bekannte sich damit unmissverständlich zu einer Programmatik der großen Zahlen: Der Name ist eine Verzerrung des Wortes Googol, welches die Zahl 10^{100} bezeichnet (vgl. Brin u. Page 1998: 1). Eine wirklich funktionale Suchmaschine, so lautet dementsprechend auch gleich die erste Feststellung des Aufsatzes, müsse über Mechanismen verfügen, die es ihr erlauben, mit dem Web zu wachsen. Als im Jahre 1994 in Gestalt des *World Wide Web Worm (WWW)* die erste Internet-Suchmaschine für den Massengebrauch online ging, indexierten die von ihr verwendeten Webcrawler (also Programme, die eigenständig Hyperlinks verfolgen und Seiteninhalte verzeichnen) etwa 110.000 Webseiten. Schon 1997 gaben die Suchmaschinenbetreiber dagegen Millionenwerte an – die an der *University of Washington* beheimatete Metasuchmaschine *WebCrawler* kannte allein zwei Millionen individuelle Seiten, das Portal *Seach Engine Watch* machte über sämtliche Suchdienste hinweg sogar rund 100 Millionen aus. Der *WWW* erhielt im Frühsommer 1994 im Schnitt bescheidene 1.500 Anfragen am Tag, während *Altavista* Ende 1997 angab, ca. 20 Millionen Suchoperationen täglich durchzuführen. Eine Suchmaschine, die angesichts dieser Entwicklung auch nur ansatzweise zukunftsfähig sein wollte, müsse laut Brin und Page dreierlei leisten: Erstens müsse sie neue Webseiten effizient verzeichnen bzw. ihre Kenntnisse über existierende regelmäßig eigenständig erneuern, zweitens müsse sie den ihr zur Verfügung stehenden Speicherplatz auf ihren Heimatservern intelligent und sparsam ausnutzen, und drittens müsse sie ihren Endnutzern ein simples, schnelles und intuitives Interface zur Verfügung stellen, welches die individuelle Suche möglichst zügig zum Erfolg führt (vgl. ebd.: 2). Im Hinblick auf das Web stellen die Autoren nämlich eine ganz ähnliche Diagnose wie Vannevar Bush und Eugene Garfield Mitte der 1940er und 60er Jahre für die Wissenschaft: Während das Informationsvolumen des Netzes seit seiner Entstehung laufend angewachsen war, sei die Fähigkeit seiner Nutzer, es zu durchsuchen und zu navigieren, konstant geblieben. Weil die existierenden Suchdienste lediglich Textsuchen betrieben, also nach bestimmten Zeichenfolgen auf Webseitentexten suchten, waren sie nicht imstande, diese in irgendeiner Form nach ihrer tatsächlichen Relevanz zu sortieren und wurden tatsächlich immer weniger nützlich, je mehr das Web anschwellte. In der Welt vor *Google* war jede Internetsuche von einer hohen Menge sogenannter *junk results* bestimmt, durch welche sich die Nutzer schlicht mit Geduld hindurchzugraben hatten, um sinnvolle Ergebnisse ausfindig zu machen. Als sinnfälligstes Armutszeugnis dieser Form von Websuche führen Brin und Page hier die Beobachtung an, dass Ende

1997 nur die vier marktführenden Suchmaschinen imstande waren, ihre eigenen Webseiten zu finden. *Google* sollte aus diesem Grunde die bloße Textsuche um ein Verfahren erweitern, das die relative Relevanz von Webseiten beurteilen und seine Treffer danach ordnen konnte (vgl. ebd.: 2f.).

Dieses Verfahren bildete das Herz der *Google*-Websuche und wurde auf den Namen *PageRank* getauft – wobei das Wort *Page* sich hier nicht etwa auf Webseiten, sondern auf die Person Larry Page bezieht. Ganz ähnlich wie der *Science Citation Index* die Verweisstruktur von wissenschaftlichen Texten auswertet, sollte *PageRank* die Relevanz von Webseiten aus ihrer Verlinkung herausrechnen. Auch hier galt es, Seiten von hoher Geltung ausfindig zu machen, die ihre Popularität an verlinkte Seiten weitergeben (vgl. ebd.: 3f.). Das Gegenstück zum *star* der Soziometrie bildet in Hypertextgefügen dabei der sogenannte *hub* – wörtlich eigentlich der ›Verkehrsumschlagplatz‹ – als eine Seite von hoher Autorität mit vielen ein- und ausgehenden Links (vgl. Mayer 2010: 71). Während in der Soziometrie und beim *SCI* jedoch ein relativ klares Verständnis davon vorausgesetzt werden konnte, was eine beliebte Person oder eine wichtige Veröffentlichung ausmacht, stand *Google* zunächst vor dem Problem, dass abhängig von den Interessenlagen seiner Nutzer die tatsächliche Relevanz einer Webseite stark variieren konnte, selbst wenn diese nach ihrem Verlinkungsgefüge eigentlich von hoher Wichtigkeit hätte sein müssen. Die funktionale Logik von *PageRank* sollte aus diesem Grund das Surfverhalten eines angenommenen Standard-Users abbilden, der mehr Flaneur als Detektiv war: *Google* sollte sich, so Brin und Page, wie ein »random surfer« betragen, der ohne ein klar gesetztes Ziel einfach Webseiten ansurft, die ihm gerade interessant erscheinen. *PageRank* erhebt entsprechend keinen Anspruch darauf, den Verlauf einer längeren Recherche oder einer ausgedehnten Auseinandersetzung mit einem bestimmten Thema antizipieren zu können. Für seine Schöpfer fungiert es vielmehr als ein auf die individuelle Suchsituation eingegrenztes »model of user behavior«, welches die unmittelbar naheliegendsten Surfentscheidungen identifiziert und anbietet (vgl. Brin & Page 1998: 4).

Über das *PageRank*-Verfahren hinaus sollte sich *Google* noch einer Anzahl anderer Tricks bedienen, um *junk*-Resultate so weit wie möglich zu unterdrücken und Seiten von hoher Relevanz an die Spitze seiner Trefferlisten zu setzen. Eines davon ist der sogenannte *anchor text*. ›Ankertexte‹ sind der Teil eines in natürlicher Sprache verfassten Webseitentextes, der als Hyperlink fungiert. Eine der Grundfunktionen von HTML ist die Möglichkeit, mittels eines einfachen Skriptbefehls eine Webadresse gewissermaßen in einer Textpassage zu ›verstecken‹, die dann einfach angeklickt werden kann und die entsprechende Webseite öffnet. Auf diese Art ist es nicht nur möglich, Hyperlinks in Webseiten einzubauen, ohne den Fließtext zu unterbrechen oder sie in einen separaten Verweisapparat auszulagern. Es kann darüber hinaus einem Link direkt eine bestimmte kulturell lesbare Bedeutung in Schriftsprache zugewiesen werden. Das komplette System der Querverweise auf *Wikipedia*

funktioniert beispielsweise nach diesem System. Dass eine solche unmittelbare Verknüpfung von *culture* und *computer layer* für ein Suchmaschinen-Ranking eine sehr aussagekräftige Quelle darstellt, ist naheliegend. Der besondere Wert, den Brin und Page ihm beimessen, gründet sich auf den Sachverhalt, dass er meist die Fremdbeschreibung einer Webseite darstellt und damit abbildet, aus welchen Interessen heraus ein Nutzer diese möglicherweise anklicken könnte (vgl. ebd.: 5). Ein weiteres Novum, mit dem *Google* gegenüber konkurrierenden Suchmaschinen auftrumpfen sollte, war die Einbeziehung geographischer Größen beim Ranking von Webseiten. Über IP-Adressen lassen sich sowohl die Standorte der Server bestimmen, auf denen Web-Inhalte gespeichert sind, als auch jene der Nutzer, die nach diesen suchen. Die räumliche Distanz kann somit ebenfalls zu einem Faktor werden, anhand dessen über die Wichtigkeit einer Webseite im Hinblick auf eine individuelle Suchanfrage entschieden werden kann (vgl. ebd.).

Wie die ihm vorausgegangenen Suchmaschinen wertet auch *Google* die Texte seiner indextierten Webseiten aus und erhebt dabei insbesondere die Häufigkeit des Auftauchens bestimmter Wörter, die in sogenannten *hit lists* verzeichnet werden. Hier unterscheiden Brin und Page allerdings abermals zwischen *plain hits* und *fancy hits*. Es kommt für die Suche nicht nur darauf an, ob ein Wort auf einer Seite auftaucht, sondern auch darauf, wie und wo es genau in Erscheinung tritt. Ein *plain hit* bedeutet lediglich, dass das Wort im Fließtext der Seite gefunden wurde. Ein *fancy hit* hingegen beschreibt einen Treffer an besonders exponierter Stelle auf der Webseite oder sogar im HTML-Code selbst. Ein *fancy hit* kann z.B. bedeuten, dass das gesuchte Wort im Titel einer Webseite oder in einer Überschrift fällt – im Gegensatz zu älteren Suchmaschinen sollte *Google* nämlich auch erkennen können, ob Textpassagen durch Fett-, Kursivsatz oder Unterstreichung hervorgehoben sind (vgl. ebd.: 5). Ein *fancy hit* wäre aber auch ein Auftauchen des Wortes in einer URL, einem Ankertext oder in mit der Webseite assoziierten Metadaten (vgl. ebd.: 9).

Kurzum: *Google*, wie Brin und Page es 1998 in Aussicht stellten, war zunächst nichts anderes als ein Bündel von Algorithmen, die Muster und Zusammenhänge erkennen und operationalisieren sollten, ohne dabei auch nur im Ansatz imstande zu sein, sie inhaltlich zu begreifen. Suchmaschinen verstehen keine kulturellen Inhalte, aber sie können unseren Umgang mit kulturellem Wissen unter ganz bestimmten Aspekten mathematisch auswerten. Die Schrift allein verbürgt – wie Laßwitz, Borges und vier Jahre vor *Google* schließlich noch Langford demonstrieren – für sich allein genommen keinen hinter den Worten wirkenden Geist, und dementsprechend führt die reine Textsuche meist nicht zum benötigten Wissen. Sie kann den Text nicht an die Welt zurückbinden, in welcher er gelesen wird. Erst im Zwischenraum der Intertextualität – dies ist eben der ›museale‹ Charakter des Netzes – werden Texte bedeutsam. *Google* kann diesen Zwischenraum nicht lesen, ihn wohl aber nach rein funktionalen Maßstäben vermessen.

Bevor diese Vermessung aber stattfinden kann, müssen die einzelnen Webseiten erst einmal erhoben werden, und nach Brin und Page ist dies der schwierigste Teil des gesamten Indexierungsprozesses. Die Webcrawler-Programme, deren Aufgabe es ja ist, systematisch so viele Links wie möglich zu verfolgen und Daten zu den so gefundenen Seiten zu erheben, müssen mit einem riesigen, nur begrenzt standardisierten digitalen Hypertextsystem umgehen, das sich der Kontrolle der Suchmaschinenbetreiber in weitesten Teilen entzieht. Sowohl die Seiten selbst als auch die Server, auf denen sie gespeichert sind, unterscheiden sich in Standards, Formaten und allgemeinem Aufbau mitunter erheblich. 1997 konnte ein Crawler ca. 300 simultane Verbindungen offenhalten und ein übliches System von vier gleichzeitig koordiniert operierenden Crawlern etwa 100 Seiten pro Sekunde indexieren. Bei derart großen Zahlen und einer solchen Geschwindigkeit liegt auf der Hand, dass eine Suchmaschine von der Größe *Googles* im laufenden Erhebungsverfahren nicht zwischen Webseiten differenzieren oder diskriminieren kann. Wer seine Seite nicht indexiert sehen möchte, der kann dies nur sicherstellen, indem er sie technisch so anlegt, dass sie von Webcrawlern nicht erfasst werden kann oder ihre Zugriffsversuche von vorneherein abgeblockt werden. Der Crawler selbst ist nicht imstande, ein sprachliches Statement auf der Webseite zu verstehen – und die Datenmenge wiederum viel zu groß, als dass sie jemals sinnvoll von Menschen gesichtet werden könnte (vgl. ebd.: 10).

Die Frage nach der Legitimität einer solch großflächigen und für die einzelnen Webangebote sowie ihre Inhalte weitgehend blinden Datenerhebung verschärfte sich im Jahre 2005, als *Google*, nunmehr bereits der unumstrittene Marktführer auf dem Gebiet der Websuchen, die Firma *Urchin Software Corp.* und mit ihr die Rechte an der Software *Urchin* aufkaufte. Bei *Urchin* – was wörtlich etwa ›Straßenkind‹ bedeutet – handelte es sich wie bei *PageRank* um ein Programm, welches das Web zwischen den eigentlichen Webseiten zu vermessen trachtet. Während *PageRank* allerdings die Verlinkung und damit gewissermaßen die Architektur des semantischen Netz-Raumes ausmisst, hatte *Urchin* nun die Aufgabe, nicht etwa die Seiten, sondern die Nutzer zu beobachten – bzw. die Fußspuren ihrer digitalen Dubletten, welche Marcos Novak mit dem Begriff des ›Navigators‹ belegt hat. *Urchin* leistet dies, indem es die Log-Dateien von Webservern ausliest, jene Dateien also, in welchen die Verbindungen verzeichnet sind, die von anderen Computern zu diesem Server aufgenommen wurden. *Urchin Software Corp.* wurde zur Keimzelle von *Googles* 2006 geformtem Firmenzweig *Google Analytics*, der heute zu den wichtigsten Pfeilern des Geschäftsmodells gehört. *Google Analytics* bietet Webseitenbetreibern die Möglichkeit, detaillierte Informationen über die Nutzer ihrer Seiten zu erhalten – im Austausch dafür, dass *Google* selbst diese Daten natürlich ebenso zur Verfügung stehen. Die einzelnen Webmaster müssen hierzu lediglich einen Account anlegen und ein paar Zeilen Javascript in den Code ihrer Webseite einfügen. *Google* kann dann jeden

Nutzer ›tracken‹, der sich auf vom Analytics-System indexierten Seiten bewegt.⁴ Im Jahre 2010 waren dies nach Angaben der *Amazon.com*-Tochter *Alexa Internet* immerhin allein die Hälfte der 1.000.000 meistbesuchten Seiten im WWW.⁵ Die Identifikation der User kann auf verschiedenem Wege erfolgen, so z.B. anhand der IP-Adresse oder durch ›Cookies‹, welche *Google Analytics* auf den Rechnern der Endnutzer ablegt. Theoretisch kann ein bestimmter Computer sogar anhand der Browsereinstellungen erkannt werden. Im Idealfall besitzt der Nutzer auch noch einen mit seinen persönlichen Daten versehenen *Googlemail*-Account. Sein komplettes Surfverhalten, alle seine Suchanfragen und alle von ihm auf *Google Search* angeklickten Treffer sind dann zentral an einer Stelle gespeichert und mit seiner individuellen Person identifizierbar⁶ – was es *Google* wiederum ermöglicht, noch genauer personalisierte Navigationshilfen im Web bereitzustellen und zukünftige Entscheidungen für und gegen das Anklicken bestimmter Links noch akkurater zu antizipieren.

In seinem Text über das virtuelle Museum als neue kulturelle Ausdrucksform identifizierte Werner Schweibenz schon 2001 – als der Siegeszug von *Google* gerade erst Fahrt aufnahm – als eine entscheidende und für jede Form von Kulturvermittlung in Netz prägende Eigenart digital-virtueller Massenmedien ein ihnen ganz eigenes Zusammenspiel von *broadcasting* und *narrowcasting*: Sie verteilen große Mengen von Information an große Massen von Menschen, sind aber gleichzeitig imstande, der Einzelperson ein sehr stark auf sie zugeschnittenes und personalisiertes Angebot aus ihrem Datenfundus zu machen (vgl. Schweibenz 2001: 12). *Google* ist für diese Tendenz innerhalb der digitalen Medienkultur und -wirtschaft nur das größte und sinnfälligste Beispiel. Nahezu alle noch existierenden Suchmaschinen arbeiten heute mit Rankingverfahren, und ihnen allen ist gemein, dass deren grundsätzliche Funktionsprinzipien zwar bekannt, die konkret zum Einsatz kommenden Algorithmen aber Betriebsgeheimnisse sind. Neben *Google Analytics* sind inzwischen noch hunderte weiterer Tracking-Dienste im Netz unterwegs, die häufig Daten für sehr spezielle und begrenzte Brancheninteressen sammeln. Und nicht nur die abstrakten Daten werden von dieser Kybernetisierung ergriffen, sondern auch konkrete Dinge. Jeder Nutzer von *Amazon* kennt jenen Startbildschirm, der ihn auf Produkte aufmerksam macht,

4 Angemerkt sei, dass der Endnutzer sich dem Tracking durchaus verweigern kann, indem er entweder ein entsprechendes Plugin in seinem Browser installiert (*Google* bietet ein eigenes an, es stehen aber auch diverse von Drittanbietern zur Verfügung) oder das Javascript insgesamt ausschaltet.

5 Vgl. <http://s3.amazonaws.com/alexa-static/top-1m.csv.zip> vom 15.05.2018.

6 *Google* bietet zwar ein Löschen dieser Daten an, dies jedoch bezieht sich nur auf das für den Benutzer sichtbare ›Logbuch‹ seiner *Google History* – *Google* selbst behält eine Logdatei, versichert jedoch, die gesammelten Daten nach 18 Monaten zu anonymisieren, vgl. <http://support.google.com/accounts/answer/54068?hl=en> vom 23.01.2015.

die dem System im Hinblick auf sein bisheriges Konsumverhalten relevant erscheinen. Diese Relevanzeinschätzung wiederum basiert logischerweise auf der Beobachtung des Kaufverhaltens der breiten Masse der Kunden.⁷ Dabei werden einzelne Artikel im Sortiment zugleich dynamisch mit anderen verlinkt, welche oft zusammen mit diesem bestellt werden. *Amazon* besitzt zwar auch ein System von Produktkategorien, mit dem sich z.B. direkt die romantischen Vampir-Romane oder die Komödien-DVDs aufrufen lassen, sein eigentliches Herz ist aber das assoziativ angelegte Navigationssystem.

Überhaupt präsentieren sich Auswertungsalgorithmen wie jene von *Google* oder die *Amazon*-Produktempfehlung zunächst einmal als logische Weiterentwicklungen des Assoziationsparadigmas, das Vannevar Bush der Memex-Maschine zugrunde gelegt hat. War es dort noch der individuelle Nutzer, der einer individuellen Maschine seine eigenen Assoziationsmuster einimpfen und sie damit außerhalb seines eigenen und notwendigerweise vergesslichen Verstandes verstetigen sollte, sind es nun *Kollektive* von Nutzern, die einer zumindest in unserem Nutzungserleben weitgehend ortlosen bzw. kaum lokalisierbaren Software-Zentralinstanz ihre Assoziationen nur mehr ganz implizit mitteilen – indem sie eben suchen, anklicken, womöglich bestellen oder herunterladen. Das ganze vernetzte Gefüge unserer kulturellen Wirklichkeit, aller unserer Zeichen, aller unserer Texte und aller ihrer Bedeutungen scheint virtuell (und vor allem: vermessbar!) verborgen zu sein in der Architektur des Hypertextes und der Art, wie individuelle Nutzer einzelne Webseiten abrufen. Der Code allein spricht nicht über die Welt, aber die Art und Weise, wie er gelesen wird, sehr wohl – und diese ist, so lautet die dem *Google*-Prinzip zugrundeliegende Annahme, eben kein aller Technik entrücktes Gemauschel im Raum des Geistes und der Emotion, sondern ein in seinem großen Umriss durchaus statistisch modellierbares Gewebe von Zugehörigkeiten und Abgeschiedenheiten, von Nähen und Distanzen, über welches letztlich nicht die Netz-Detektive, sondern die Netz-Flaneure abstimmen. Dies tun sie mit den metaphorischen Füßen: Das Besondere am Hypertext ist eben, dass der Text nicht mehr nur die Entscheidungen des Autors abbildet, sondern zugleich jene des Lesers einfordert, und jede dieser Leserentscheidungen muss dem System in Form einer klaren Auswahl dieses oder jenes Hyperlinks mitgeteilt werden, womit sie unweigerlich serverseitig erfassbar wird. Das Web steht nicht immer und überall im direkten ›Dialog‹ mit seinen Nutzern, aber wir erklären ihm grundsätzlich mit jedem Klick aufs Neue unsere Absichten.

7 *Amazon.com* nutzt dabei ein kompliziertes System, das ›Cluster‹ von Zugriffen auf Artikel auswertet und hierbei sowohl Nutzer- als auch Artikelgruppen einbezieht (vgl. Linden, Smith u. York 2003: 76f.; 78f.).

4.2.3 Das gezähmte Netz: Vom Flanieren zum Finden

Der Erfolg sowohl des Hypertext-Formates als auch jener des Nutzer-Trackings ist sicherlich auch siebzig Jahre nach *As We May Think* noch im Kontext einer evidenten Unfähigkeit klassischer Wissensordnungs- und Abrufsysteme zu lesen, mit rasanten Informationsexplosionen mitzuhalten. Die institutionelle und personelle Autorität, die sich mit der Idee der Autorschaft verbindet und die in einem Bibliothekskatalog ebenso präsent ist wie in einer Museumsausstellung, hat grundsätzlich nur eine begrenzte Reichweite, die sich aus den Einschränkungen des menschlichen Organismus ergibt: Eine Gruppe von Bibliotheksmitarbeitern kann nur eine begrenzte Zahl von Büchern bibliographisch erschließen, eine Gruppe Kuratoren nur eine bestimmte Menge von Exponaten verwalten und eine noch kleinere in die Ausstellung aufnehmen. Dabei sind solche Expertensysteme, wie Konrad Becker feststellt, stets befangen in ihren jeweiligen disziplinären Vorurteilen, die nicht immer das Ergebnis kritischer Reflexion und Auseinandersetzung sind, sondern allzu oft quasi-religiöse Züge aufweisen:

Professionelle Kategorisierungsexperten bemühen sich, das Kontextabhängige und Temporäre um jeden Preis zu meiden, und enden doch immer mittendrin. [...] Fast zwangsläufig dominieren eigene Interessen und Anforderungen der Katalogproduzenten über objektivere Bedürfnisse von Navigation in komplexen Welten. Sie züchten kognitive Verwaltungstechnologien, die kulturelle und subjektive Vieldeutigkeit und das Schillern kontextabhängiger Aussagen nicht erkennen. Vorstellungen einer objektiven Ordnung des abstrakten Raumes entstehen auf Grundlage religiöser Ideen von makelloser Reinheit. Solche Illusionen werden von gefährlichen Ideologien kybernetischer Kontrolle genährt und vom Glauben, dass die manifeste Welt auf einen einzigen Standpunkt zu reduzieren wäre. (Becker 2010: 185)

Damit möchte Becker aber keineswegs ein Loblied auf die völlige Kontextabhängigkeit anstimmen: Vielmehr, so seine These, gibt es derzeit schlicht keine ideale Form der Zurechtfindung in komplexen Wissensräumen von der Größe des WWW. Zwar sei eine kategoriale Ordnung seiner Inhalte weder machbar noch für die Masse der Nutzer in irgendeiner Form zielführend, zugleich jedoch vermittele eine rein assoziative Zugriffsform allzu leicht falsche Vorstellungen von Verfügbarkeit und schaffe sich darüber hinaus immer auch ihre ganz eigenen Irrwege:

Der Versuch der elektronischen Aufbereitung von Informationen und Ressourcen kann durch überkommene Gewohnheiten und veraltete Strategien aus vorangegangenen Ansätzen der Strukturierung von Wissen sehr unzureichend ausfallen. Digitale Information braucht kein Regal, und es stellt sich die Frage, inwieweit vordefinierte Kategorisierung überhaupt eine gute Idee ist. Ein Hauptgrund für den Erfolg von Google war das Fehlen virtueller Regale, von im

Voraus konstruierten und immer auch eigenartigen Dateistrukturen. Aber bei Regalflächen, auch wenn sie den Wissensraum seltsam verzerren, ist zumindest einfach zu sehen, ob sie voll oder leer sind. (Ebd.)

Das Heimtückische an der Assoziation im näherungsweise Absoluten ist also, dass – ganz ähnlich wie die Textsuche in Langfords *Net of Babel* immer zum Erfolg führen wird – uns die Assoziation in einem ausreichend großen und vernetzten Hypertextsystem letztlich fast immer in irgendeiner Form zum nächsten Textsegment führen wird. Ein schwarzer Bildschirm, der uns sagt ›Du, Nutzer, hast die Grenze sinnvoller Assozierbarkeit erreicht, ab hier gibt es nichts mehr zu lernen‹ ist in der Konzeption des Webs nicht vorgesehen. Während die Bibliothek von Babel irgendwann durchschaut und der Umstand, dass ein Text in überhaupt keinem kulturellen Bezug zur Außenwelt steht, als Normalzustand erkannt wird, lädt das Netz uns immer wieder zum abduktiven Trugschluss ein: Die statistische Kontingenz der Zugriffsmuster erzeugt eine kulturell-narrative Kontingenz der Sinnzusammenhänge, aber diese entsteht selbst nicht in der Domäne des Kulturellen. Bevor unsere kulturellen Schemata das digitale Objekt nach Hui aus den Relationen zwischen digitalen Informationspartikeln entstehen lassen können, nehmen solche Navigationswerkzeuge die gemittelten Mechanismen unserer kollektiven Kognition vorweg und präsentieren uns das, was sie als wahrscheinlich kommensurabel für uns identifiziert haben.

Die Vorstellung vom völlig emergenten und gänzlich situativ um unsere eigenen Interessenspfade herum entstehenden Cyberspace ist, wie 2012 der weißrussische Publizist Evgeny Morozov feststellt, in erster Linie Leitbild und frommer Wunsch der Advokaten einer emanzipatorischen Netzkultur, die möglicherweise längst technisch unmöglich geworden ist. Morozov sieht darin keinen Sieg des Flaneur-Paradigmas über jenes des kritischen Detektivs – sein Text *The Death of the Cyberflâneur* sieht den Flaneur (offensichtlich) vielmehr als das erste und bedeutendste Opfer eines nach statistisch ausgewertetem Massenverhalten geordneten Webs (vgl. Morozov 2012).

Der Flaneur nämlich, so Morozov, ist originär eine Figur, die sich zwar mit großer Selbstverständlichkeit in der Masse bewegt, sich aber zu keiner Zeit von ihr vereinnehmen lässt. Sein Begriff vom Flaneur im Datenraum ist also sehr viel näher jenem Lev Manovichs verbunden als jenem Roberto Simanowskis. Der Flaneur nach Morozov zeichnet sich dadurch aus, dass er zwar gern durch die Einkaufsstraße tändelt, niemals jedoch der Versuchung des unreflektierten Konsums erliegt. Der Flaneur im Baudelaire'schen Sinne wolle am städtischen Leben nicht funktional teilnehmen, sondern schauen, beobachten, verinnerlichen und – selbstverständlich – darüber erzählen und schreiben. Die Stadt war für ihn also eher Anschauungs- als Handlungsraum. Der Flaneur sei daher immer Kulturkritiker – allerdings einer, der sich nicht etwa vom Feldherrenhügel einer wie auch immer definierten Hochkultur herab äußert, sondern aus dem kulturellen Milieu der Straße heraus über die Straße und die

sich in ihr bewegenden (und für sich selbst blinden) Massen schreibt. Dieses Bild vom Internetnutzer als einem lustwandlerischen, zugleich aber auch nachdenklichen und kritischen ›Entdecker‹ schlägt sich nach Morozovs Lesart schon in den Namen der ersten Browserssoftwares nieder: Bezeichnungen wie *Internet Explorer* und *Netscape Navigator* waren in der Frühzeit des WWW attraktive und affirmative Metaphern für eine mündige Internetnutzung (vgl. ebd.).

Von diesem Nutzerbild zieht Morozov eine interessante Parallele zum Schicksal des historischen (Straßen-)Flaneurs: Dieser nämlich sei an einer Veränderung seiner Umwelt zugrunde gegangen, die er nicht zu kompensieren imstande war. Beginnend im Frankreich des Louis Napoleon und seines Pariser Stadtplaners Georges-Eugène Baron Haussmann hätten Neuordnungen des urbanen Raumes ihm seine wichtigste Lebensgrundlage entzogen: die verwinkelten, labyrinthischen Seitenstraßen, die das Mittelalter den europäischen Metropolen hinterlassen hatte und die dem Flaneur sowohl das äußere Verlorengehen im öffentlichen Raum, als auch das innere im Privaten der Gedanken ermöglicht hatten. An ihre Stelle traten weite, begradigte Boulevards, die nicht nur dem Fußgänger, sondern wenige Jahrzehnte später gerade auch dem Automobil ein schnelles und zielführendes Vorwärtsskommen ermöglichen sollten, dabei aber zugleich leichter sauber zu halten und – nach 1848 nicht mehr nur in Frankreich ein entscheidender Faktor – schwerer zu verbarrikadieren waren. Mit den Boulevards kam die Gasbeleuchtung, welche die Stadt zunehmend auch als Raum der abendlichen Freizeitgestaltung interessant werden ließ. Die Nacht gehörte nicht länger Flaneuren und Halbweltlern, sondern wurde vom Bürgertum und seinen Zerstreuungen erobert. Viele klassische Einkaufsstraßen mit kleinen Einzelgeschäften fielen der Konkurrenz kompakter Kaufhäuser zum Opfer, in denen nicht mehr gebummelt, sondern zügig, effizient und verkaufsorientiert bedient wurde. Und nicht zuletzt wurde das Flanieren immer gefährlicher: Der zunehmende (und vor allem auch: zunehmend motorisierte) Verkehr machte das versonnene Spazieren zu einem Unterfangen, das durchaus im Krankenhaus oder auf dem Friedhof enden konnte. Nach Morozov waren es eben diese Entwicklungen im räumlichen Gefüge der Stadt, die schließlich auch die größten Flaneure – Proust, Baudelaire, Sainte-Beuve – ins innere Exil ihrer Wohnungen trieben (vgl. ebd.).

Im World Wide Web der Gegenwart sieht Morozov ganz ähnliche Mechanismen am Werk wie in der zweiten Hälfte des 19. Jahrhunderts in den Großstädten Europas: Der Cyberflaneur existierte, so seine These, in einem Web, das noch keine (oder nur relativ ineffiziente) Suchmaschinen kannte und noch nicht durchgreifend von kommerziellen Akteuren erobert worden war. Das Netz-Äquivalent zu den mittelalterlichen Sträßchen und Gassen des städtebaulich noch nicht modernisierten Paris sieht er dabei in den »Arkaden« der bunten *Geocities*-Seiten der frühen 1990er Jahre, die inzwischen griffigeren Interfaces und Abrufmechanismen gewichen seien. Im Netz gäbe es keinen Raum mehr zum Flanieren, weil sich alles viel zu leicht finden lasse (vgl. ebd.).

Den Baron Haussmann des Netzes erkennt er dabei interessanterweise aber nicht etwa in Sergey Brin und Larry Page, sondern im *Facebook*-Gründer Mark Zuckerberg. Nicht nur die bis zur Trivialität vereinfachte Auffindbarkeit von Daten sei der Todesstoß für den Cyberflaneur, sondern vielmehr auch eine »Tyrannei des Sozialen« (ebd.), die auf das Herz des flanierenden Erfahrungsmodus ziele. Dienste wie *Facebook* wollen, so schreibt Morozov, jede Erfahrung im Netz zu einer sozialen machen und damit jene Anonymität und Ambivalenz vernichten, auf die der Flaneur existenziell angewiesen ist. Die Einsamkeit, die in definiert, soll schlicht nicht mehr möglich sein.

Indes sind die hinter der Suchmaschine *Google* und dem *social network Facebook* stehenden Absichten und Geschäftsmodelle natürlich ohnehin ganz ähnliche: Ihre Ware und Währung sind die Nutzungsprofile derer, die ihnen mehr oder weniger bewusst und willig Informationen über sich selbst zur Verfügung stellen und damit eine Optimierung von Online-Angeboten auf ihre Interessenlagen hin ermöglichen. Dass hier natürlich der Weg für allerhand selbsterfüllende Prophezeiungen geebnet ist, muss kaum hervorgehoben werden. *Google* kann auf die Bedürfnisse und die Lebensweise des Flaneurs gar keine Rücksicht nehmen, weil seine epistemischen Spaziergänge statistisch kaum zu erfassen sind – um seinen assoziativen Gedankenketten zu folgen, müsste man sie eben inhaltlich verstehen. *Google* kann einzig die Bewegungen von Massen verfolgen und anhand derer zwar nicht die Natur von *connectedness* vermessen, wohl aber ein mathematisches Modell ihrer relativen Stärke zeichnen. Weil *Google* aber nicht nur passiv beobachtet, sondern aufgrund seiner Auswertungen die »Relevanz« von Wissensinhalten unmittelbar an seine Nutzer zurückspielt, gewichtet die Software die Autorität bestimmter Webseiten unweigerlich immer weiter zu einem bestimmten »Gondelende« (vgl. Jeanneney 2007: 44) hin: Eine Seite steht hoch oben in der *Google*-Trefferliste, weil sie oft angeklickt wird – und sie wird oft angeklickt, weil sie hoch oben in der *Google*-Trefferliste steht. Der Historiker Jean Noël Jeanneney spricht hier im Jahre 2007 – unmittelbar nach dem Ende seiner Dienstzeit als Direktor der französischen Nationalbibliothek – unter dem Eindruck von *Googles* großangelegtem Bücherdigitalisierungsprojekt *Google Books* von einer »natürlichen Kybernetik« (ebd.), die nunmehr mit der Suchmaschine ihre technische Untermauerung und Vollendung erfahren habe: Kulturelle Wissensbestände neigen zur rekursiven Selbstbestärkung und damit zur Privilegierung von Wissen, das sich leicht an etablierte Kenntnisschätze anschließt. Diesen »epistemischen Selbstläufern«⁸ gegenüber stehen unweigerlich ausgedehnte Landschaften »schwierigen« Wissens, das nicht den Majoritäts-Lesarten entspricht und am fernen Ende der Gondel marginalisiert wird (vgl. Jeanneney 2007: 44). Die Verfolgung solch randständigen

8 Stefan Rieger entwickelt den Begriff des »epistemischen Selbstläufers« bezeichnenderweise an den technischen Apparaturen physiologischer Experimentalsysteme des 19. Jahrhunderts (vgl. Rieger 2012).

Spezialwissens ist das Privileg individualisierter Rezeption – der Flaneur lässt es lustwandelnd auf sich zukommen, der Detektiv stellt ihm angestrengt nach, aber beide Figuren scheinen keinen rechten Platz mehr zu haben in einer Informationswelt, die auf schnellste Verfügbarkeit und umweglosen Abruf ausgerichtet ist.

Jeanneney sorgt sich mit Blick auf *Google Books* ganz konkret um die Geltung der frankophonen Literatur in der globalisierten Welt. Weil *Google* erstens zugleich das Mess- und Etablierungswerkzeug einer hierarchischen Ordnung unter kulturellen Inhalten sei und zweitens die Popularität des Populären bestärke, begünstige sein Ausgreifen in die Welt der gedruckten Bücher im kolossalen Maße die Verbreitung englischsprachiger Texte, und damit letztlich auch die der englischen Sprache schlechthin (vgl. ebd.: 6ff.). Da Suchmaschinen durch geschickte mathematische Tricks die Grenze zwischen *computer* und *culture layer* überbrücken, müssen sich ihre eigentlich völlig kulturfremden Algorithmen unweigerlich Debatten über ihre kulturellen Implikationen stellen – und damit auch über ihre Rolle in Auseinandersetzungen zwischen privilegiertem und marginalisiertem Wissen sowie den Identitäten und Lebensentwürfen, die sich mit diesem verbinden. Letztlich aus reiner Mathematik bestehende Software wird zum Abbild von Ideologien und damit implizit politisch: In einer Welt, in der Information längst überwiegend digital gespeichert, verbreitet und abgerufen wird, programmieren Programmierer eben nicht mehr nur Computer, sondern auch die Kulturen, die sich ihrer bedienen. Damit stehen etablierte Kulturinstitutionen vor ganz neuen Herausforderungen – und sehen sich mit ganz neuer Konkurrenz konfrontiert.

4.3 ANDRÉ MALRAUX: DAS IMAGINÄRE MUSEUM

Die Geschichte des Museums ist, wie das erste Kapitel dieser Arbeit abrissartig darstellen konnte, eng verwoben mit gesellschaftspolitischen Visionen und Utopien, die sich durchaus als *Programme* begreifen lassen: *Vor-Schriften* des Denkens und der Welterfahrung, die sowohl deren *Inhalt* als auch ihren *Modus* betreffen und die das Museum bei seinen Besuchern anregen, kultivieren, verfestigen soll. Von der Vermittlung einer vermeintlich göttlichen Ordnung in den Wunderkammern der Frühneuzeit über die Humboldt'schen Musentempel zu den »emanzipatorischen« Ausstellungskonzepten der Nachkriegszeit wollten Museen nicht nur ein spezifisches Wissen, sondern auch bestimmte Formen und Philosophien des Lernens transportieren. Eine Kontinuität der Museumsgeschichte, die sie nun mit jener der digitalen Medien verbindet, war dabei das Netzwerk als funktionales Leitprinzip. Eine andere, die immer noch als Barriere zwischen dem Museum und dem Web steht, ist die Innen-/Außen-Trennung und mit ihr die naturgemäße Beschränktheit aller Sammlungen und