

Big Data Is Watching You

Digitale Entmündigung am Beispiel von Facebook und Google

Rainer Mühlhoff

1 Einleitung

Die Debatte um Datenschutz im Internet, die seit den Enthüllungen durch Edward Snowden 2013 vermehrt in der medialen Öffentlichkeit geführt wird, ist weitgehend auf das Gefahrenszenario eines unbefugten Zugriffs auf eigentlich private Daten durch Dritte, etwa durch Geheimdienste, staatliche Akteure, Hacker oder Kriminelle, fokussiert. Diskutierte Szenarien sind der Einbruch in einen Server, bei dem etwa Kundendaten gestohlen werden; das Abhören von Telekommunikation durch Geheimdienste, Polizei und Verfassungsschutzorgane; oder Attacken durch Viren, Malware und Phishing-Techniken gegen ahnungslose Nutzer_innen, etwa mit dem Ziel, an ihre Online-Banking-Daten zu gelangen oder ihre E-Mail-Accounts zu übernehmen.

Zur gleichen Zeit sind Unternehmen wie Google, Apple, Facebook oder Amazon (fortan gemeinsam abgekürzt »GAFA« genannt) als umfassende Datensammler bekannt. Im Rahmen ihrer Services erfassen, verarbeiten und vermarkten sie Personen- und Nutzungsdaten – dies in einem Umfang und mit einer strukturierten Informationsauflösung, die die Datenvorräte von staatlichen Stellen und intrusiven Datensammlern weit übersteigt. Gemessen hieran fällt die öffentliche und politische Problematisierung der Datenerhebungspraxis dieser Unternehmen gering aus. Selbst wenn kritisch diskutiert wird, dass diese Unternehmen Daten sammeln, verdeckt ein oft alarmistischer Ton den eigentlich entscheidenden Punkt, nämlich die Schaffung eines kollektiven Bewusstseins dafür, *wie* diese Unternehmen ihre Daten sammeln. Eine Beschwichtigung erfährt die Debatte überdies oft schnell durch den Hinweis, dass dieser Fall auch rechtlich anders gelagert sei als das Problem staatlicher Überwachung oder intrusiver Da-

Mühlhoff, R.; Breljak, A.; Slaby, J. (Hg.): *Affekt Macht Netz. Auf dem Weg zu einer Sozialtheorie der Digitalen Gesellschaft*. transcript 2019, S. 81-107. DOI: 10.14361/9783839444394-004.

tenbeschaffung, da die Nutzer_innen ihre Daten auf den privaten Plattformen von GAFA meist freiwillig, wissentlich, unter Einwilligung in Nutzungsbedingungen und somit in bewusstem Verzicht auf bestimmte Rechte preisgeben.

Die Plattformbetreiber ihrerseits betonen, dass sie alle Nutzerdaten gemäß der gängigen Gesetzgebung behandeln und den Nutzer_innen umfangreiche Einstellungsmöglichkeiten zur Kontrolle der öffentlichen Sichtbarkeit ihrer Daten (siehe Facebook) und zur »Absicherung« ihrer Accounts anbieten.¹ Diese Praxis lässt GAFA als Plattformbetreiber tendenziell sogar als verdeckte Gewinner aus der Verunsicherung über Datensicherheit post-Snowden hervorgehen. Denn sie lenkt von dem Problem ab, das ihr eigenes massenweises Aufzeichnen von Nutzungsdaten darstellt, indem sie das Bild erzeugt, die Nutzer_innen könnten freiwillig und selbstbestimmt den Umgang ihrer auf der Plattform hinterlegten Daten kontrollieren.²

Um was für eine Form der ›Freiwilligkeit‹ auf Seiten der Nutzenden handelt es sich bei der Erfassung von Nutzerdaten, die speziell *im* Rahmen von Nutzungsbedingungen und selbst gestaltbaren »Sichtbarkeiteinstellungen« erfolgen? Es macht sich kaum jemand Illusionen darüber, dass die Nutzungsbedingungen oder Datenschutzhinweise, die hier und da aufpoppen, überhaupt von Vielen gelesen werden. Dies nicht nur, weil es Zeit kostet, im falschen Moment daher kommt oder die seitenlangen juristischen Klauseln ein Gefühl der Ohnmacht erzeugen, sondern auch weil es gar keine Möglichkeit gibt, ihnen zu widersprechen – es sei denn, man verzichtet gleich ganz auf die Benutzung des entsprechenden Services. Gegenüber GAFA besteht so etwas wie ein fatales Ausgeliefertsein an die Zwickmühle subjektiv empfundener Unverzichtbarkeit dieser Services, denen für eine Mehrheit der Nutzer_innen eher der Status einer Infrastruktur denn einer Dienstleistung zukommt.

1 | Mit seitengroßen Anzeigen im Stil des Native Advertisings versuchte Google sich zum Beispiel im April 2017 über große deutsche und europäische Zeitungen als verantwortungsbewusster Akteur beim Thema Datenschutz darzustellen. So brachten die *Süddeutsche Zeitung*, *die Zeit* und *Spiegel Online* am 7. April 2017 gesponserte Inhalte der Firma Google im Stil eines Zeitungsartikels mit dem Titel »Was macht ihr eigentlich mit unseren Daten?«, Dachzeile »Rede und Antwort«. Ähnliche Aktionen gab auch Facebook in Auftrag, etwa eine Kampagne im Februar 2018 unter der Überschrift »Du hast die Kontrolle über deine Daten auf Facebook«, die europaweit in Magazinen und Tageszeitungen gedruckt wurde.

2 | Der Diskurs um Datenschutz auf Plattformen wie zum Beispiel Facebook entwickelt sich also in die Richtung, dass Facebook feinschrittige Einstellungsmöglichkeiten für die »Sichtbarkeit« einzelner Nutzerdaten für andere bietet und damit das Gefühl vermittelt, die Fragen von Datensicherheit und Privatsphäre ernst zu nehmen. Dennoch besitzt die Plattform selbst alle diese Daten und wertet sie auch aus, denn für die Generierung *abgeleiteter* Daten, etwa zu Risk-Controlling und Marketingzwecken, ist keine Weitergabe der persönlichen Daten selbst erforderlich.

Dieses Ausgeliefertsein dokumentiert sich auch in einem kollektiven Verdrängungs- und Herunterspielungseffekt in Bezug auf die persönlichen, sozialen und gesellschaftlichen Folgen der Nutzung durch die dabei anfallenden Massendaten. Er kommt in einem breiten Spektrum psychologischer Einstellungen und Bewältigungsmuster zum Ausdruck: vom fatalistischen Achselzucken, dem naiven Glauben »Ich habe doch nichts zu verbergen« und dem resignativen »Über mich haben die schon so viele Daten, da machen die paar mehr jetzt auch nichts mehr aus« bis zu Haltungen der begeisterten Affirmation, welche gegenüber den Nachteilen eher die Freiheits- und Entfaltungsmöglichkeiten, die ökonomischen und politischen (Gewinn-)Chancen oder die lebensstilistische Überlegenheit der neuen Technologien in den Vordergrund stellen.

In dieser Situation zwar divergierender, insgesamt jedoch herunterspielender oder verdrängender Umgangsweisen ist zu beobachten, wie sich in der Debatte um Datenschutz gesamtgesellschaftlich ein impliziter *liberalistischer* Konsens über die Trennung von Staat und Privatökonomie etabliert hat, der auch in weiten Teilen politisch linker Kreise geteilt wird: Der Diskurs um Überwachung, Datensicherheit und Privatsphäre reproduziert die Gegenüberstellung von Staat und Privatökonomie, die im Kern liberaler Gesellschaftssysteme steht. Gegenüber dem Staat herrscht (mitunter zurecht) ein grundsätzliches Misstrauen in Bezug auf Datenerhebung, das jedoch gegenüber ökonomischen Akteuren effektiv nicht oder nicht in gleicher Form verbreitet ist. In der Situation fatalen Ausgeliefertseins an Plattformunternehmen lässt es sich hier und da fast als eine psychologische Verschiebung lesen, sich über Geheimdienste zu empören, deren verdeckte Vorgehensweise an einen gewaltvollen und repressiven Staats- und Polizeiapparat gekoppelt und deshalb vergleichsweise direkt kritisierbar ist, während sich im Hinblick auf die eigene Benutzung kommerzieller Plattformen die Haltung durchsetzt, dass ja jede_r selbst entscheide und selbst kontrolliere, ob und was für Informationen über sich oder über andere man diesen Plattformen übermittelt. Man stellt Daten zur Verfügung, aber eben *freiwillig* und *wissentlich*, weil es einem egal ist oder weil man es will, und weil man glaubt, dass die Datenpreisgabe ja nur einen selbst betrifft.

In diesem Essay möchte ich diesen Aspekt der »Freiwilligkeit« und »Wissentlichkeit« näher beleuchten und anhand technischer Beispiele in Frage stellen. Nach der Methode einer neuen Technik-Phänomenologie werde ich die konkrete Gestaltung von Mensch-Maschine-Interaktion als Kontextfaktor »freiwilliger« Nutzungsdatenerhebungen im Internet untersuchen. Daran wird sich zeigen, dass das Design von Interfaces und Benutzeroberflächen in vielen Fällen Nutzer_innen entmündigt und auf eine bestimmte reflexive Beziehung zu und Wahrnehmungsweise von technischen Artefakten hin ausrichtet. Dadurch wird die Interaktion mit dem Interface subtil so gestaltet, dass Nutzer_innen mit höherer Wahrscheinlichkeit möglichst viele Daten über sich preisgeben.

2 Fallstudien: Technologische Spielarten von »Freiwilligkeit«

Die These, die im Folgenden an drei Beispielen zu illustrieren ist, lautet: Die Erhebung von Personen- und Nutzungsdaten auf webbasierten Plattformen wie Google und Facebook erfolgt durch Techniken des Nudgings und der *verdeckten* Erhebung von Bewegungsdaten. Diese Techniken sind Gegenstand eines aktuell virulenten und von viel Kapital gestützten Diskurses zwischen Designer_innen, Programmierer_innen, Unternehmen und Technologievisionär_innen, der sich in Felder wie User Experience Design (UX Design), Search Engine Optimization (SEO) und Web Analytics hinein verzweigt. Insgesamt operiert das mediale Dispositiv, in dem die massenweise Erhebung von Personen- und Nutzungsdaten möglich ist, über eine psychologische Dimension des Nudgings und Social Engineerings, und über eine technische Dimension des Trackings.

Die psychologische Dimension – Nudging³ und Social Engineering – bezeichnet ein Feld von Techniken, die darauf zielen, potenzielle Nutzer_innen für eine Internetanwendung zu gewinnen (»onboarding«) oder innerhalb einer Anwendung zu bestimmten Entscheidungen zu bewegen, zum Beispiel etwas zu kaufen, bestimmte Daten preiszugeben oder Zugriffsrechte zu erteilen. Um Social Engineering handelt es sich dabei, wenn ein bestimmter Service auf dem Wege lebensweltlicher Verankerung und sozialer Druck- und Zugehörigkeitsmechanismen verbreitet wird, etwa im viralen Marketing oder durch die Ausnutzung von Netzwerkeffekten bei der Verbreitung von Messenger-Diensten und sozialen Netzwerkplattformen. Nudging hingegen betrifft den situativen Aspekt des Designs von Interfaces, Benutzeroberflächen, Dialogboxen, »User Experience«. Dabei ist die zentrale Frage, die in einem Diskurs zwischen Technik und Psychologie behandelt wird, wie die Ansprache eines Subjekts durch die mediale Oberfläche gestaltet werden muss, um sie auf ein bestimmtes *wahrscheinliches* Benutzungsverhalten hin zu optimieren, zum Beispiel darauf, dass die Nutzer_in einen Account anlegt, etwas kauft, der Übermittlung ihrer Daten zustimmt etc. Die zweite, technische Dimension, die mit der psychologischen jedoch verschränkt ist, betrifft Techniken des Trackings und des Quantifizierens von Benutzerflüssen und »click streams« im Netz. Es handelt sich hier um Techniken, die dezentral und unter der Oberfläche von mehr als zwei Drittel aller Websites im Internet arbeiten, um das Nutzerverhalten zu vermessen und so die empirische Grundlage zur Verifizierung von Nudging- und Social Engineering-Techniken zu liefern.⁴

3 | Vgl. für die ursprüngliche Verwendung des Begriffs »nudge« für die Idee eines »libertären Paternalismus« in Verhaltensökonomie und Public-Policy-Diskursen Thaler und Sunstein 2008 und zur Übertragung auf digitale Interfaces Mühlhoff 2018.

4 | Nach einer Statistik des *World Wide Web Consortium* setzten im Jahr 2018 rund 65% aller Websites Analytics-Tools ein, welche die Bewegungen von Nutzern auf Websites aufzeichnen. Rund 35% aller Websites verwenden zudem eigene Session-Cookies, um Nutzer_innen über ein-

Im Folgenden soll also gezeigt werden, dass die Datenerhebung auf Plattformen wie Google und Facebook auf Techniken beruht, die kategorial anders operieren als Geheimdienstspionage oder intrusiver Datenklau. Es handelt sich um Techniken, die erstens auf dem unwissentlichen Mitwirken von Nutzer_innen beruhen und die zweitens dezentral implementiert werden – also über ein komplexes Zusammenspiel von ökonomischen, subjektiven und technischen Strukturen operieren. Ich unterscheide im Folgenden drei Formen von Unfreiwilligkeit, die jeweils in den drei konkreten Fallstudien illustriert werden:

1. *Unbemerkte*, aber nicht heimliche Erhebung von Daten, zum Beispiel in der dezentralen Struktur des Klick-Trackings in der Google-Suche.
2. Subjektiv *freiwillige* und *wissentliche*, aber nicht voll informierte Weitergabe von Daten in situativen Interface Nudges. Dies wird am Beispiel der Single Sign-on Services »Google Sign-In« und »Facebook Login« illustriert.
3. Eine im vollumfänglichen Sinne *wissentliche* Erhebung von Daten, bei der die Nutzer_in aber unfreiwillig auch noch für die Arbeit der maschinenlesbaren Aufbereitung und Verwertbarmachung ihrer Daten eingespannt wird. Beispiel: Der persönliche Steckbrief auf einem Facebook-Profil.

Beispiel 1: Klick-Tracking in der Google-Suche

Die Suchmaschine Google ist aus heutiger Sicht das größte Quasi-Infrastrukturprojekt, welches aus der zweiten Welle der New Economy nach dem Platzen der Dotcom-Blase Anfang der 2000er Jahre hervorgegangen ist. Die ersten Internet-suchmaschinen der 1990er Jahre, darunter Yahoo und Lycos, arbeiteten noch wie erweiterte Telefonverzeichnisse, in denen Websitebetreiber_innen ihre Seiten aktiv mit Stichworten und kurzen Selbstbeschreibungen hinterlegen mussten, so dass eine Suchanfrage eher dem Blick in ein kommerzielles Verzeichnis wie die »Gelben Seiten« entsprach. Einen zweiten wesentlichen Schritt in der Entwicklung von Suchmaschinentechnologien bildete das Modell, automatisiert die netzwerkförmige Hyperlink-Verweisstruktur des Internet zu nutzen und von einem Einstiegspunkt ausgehend rekursiv allen Hyperlinks zu anderen Seiten zu folgen, um so das gesamte Netz für Suchzwecke zu indizieren (Prinzip des Webcrawlers).⁵ Während hiermit die manuelle Registrierung jeder einzelnen Website

zelne Besuche hinweg wiedererkennen und einer übergeordneten Session zuordnen zu können. Siehe https://w3techs.com/technologies/history_overview/traffic_analysis/all und <https://w3techs.com/technologies/details/ce-cookies/all/all>, abgerufen am 1.4.2019.

5 | Webcrawler wurden ursprünglich auch unter der Bezeichnung »Wanderer« bekannt. Der historische Prototyp ist der *World Wide Wanderer*, der 1993 von Matthew Gray entwickelt wurde und ursprünglich den Zweck hatte, das Wachstum des Internets messen zu können (Sajja und Akerkar 2012: 85 ff.).

hinfällig wurde, stellte sich jedoch weiterhin – wie bei jeder Suchtechnologie – das Problem der Gewichtung und der linearen Anordnung von Suchresultaten, um möglichst einschlägige Treffer zuerst, weniger einschlägige nachgelagert anzeigen zu können. Wichtige Kriterien, nach denen diese Gewichtung in der zweiten Generation von Suchtechnologien vorgenommen wurde, waren zum Beispiel die Häufigkeit, mit dem ein Suchstichwort auf einer Seite vorkam, oder die räumliche Nähe, in der die gesuchten Stichworte auf einer Seite anzutreffen waren, sowie die Anzahl anderer Seiten, die auf eine Seite qua Hyperlink verwiesen. Hinzu kam die Möglichkeit, die die meisten kommerziellen Suchmaschinen schon immer geboten haben, für eine Höherbewertung der eigenen Seite in den Suchresultaten etwas zu bezahlen.

Für das zentrale Problem der Gewichtung von Suchresultaten hat nun Google – und das markiert zugleich den Übergang zu einer dritten und aktuellen Generation von Suchmaschinentechnologien im Web 2.0 – eine besonders effiziente Lösung gefunden, die auf der unbemerkten Mitarbeit der Nutzer_innen basiert. Jede Person, die im Netz etwas sucht und eventuell mehr oder weniger fündig wird, ist prinzipiell eine »kognitive Ressource«, die man zur Qualitäts- und Relevanzbemessung von Suchresultaten einspannen kann. Dazu hat Google eine Infrastruktur geschaffen, die zu erfassen erlaubt, welche der auf der Resultatseite (fortan im Jargon kurz SERP, *Search Engine Result Page* genannt) gelisteten Suchresultate tatsächlich angeklickt werden und ob diese Resultate zufriedenstellend waren. Ruft man zum Beispiel in einem Firefox-Browser google.de auf und startet eine neue Suche nach dem Stichwort »aktuelle Nachrichten«, so sieht man etwa »tagesschau.de«, »n-tv.de«, »bild.de«, »fnp.de« als erste Resultate. Führt man mit der Maus über den ersten Link, der zu tagesschau.de führt, sieht man in der Statuszeile des Browsers »<http://www.tagesschau.de/>« – das suggeriert, dass der Link wirklich (direkt) zur Tagesschau führt. Schaut man in den HTML-Quelltext der Resultatseite,⁶ dann sieht dieser Link folgendermaßen aus:

```
<a href="http://www.tagesschau.de/"
  onmousedown="return rwt(this,'','','','1',
    'A0vVaw1_8NtrWpEbLefVcTd78eDR','','0ahUKEwj5t3t0YLZAhB
    ULlSwKHb-eAYcQFggUAAA','','','event')">Aktuelle Nachrichten -
  Inland Ausland Wirtschaft Kultur Sport - ARD ...</a>7
```

[Titel] ist die Standard-HTML-Struktur für einen Hyperlink, bei dem der User auf [Titel] klickt und zur Seite [Zieladresse] gelangt. Tatsächlich scheint der oben untersuchte HTML-Link also direkt zu <http://www.tagesschau.de/> zu führen (href-Parameter). Aber der Link wird

6 | In Firefox klicke man mit der rechten Maustaste in das Fenster, dann »View Page Source«. In dem unübersichtlichen Quelltext suche man nach »tagesschau.de«, um zur relevanten Stelle zu gelangen.

7 | Hervorhebungen und Zeilenumbrüche in allen Code-Blöcken vom Verfasser eingefügt.

vom Browser erst geöffnet, wenn man beim Anklicken den Mausbutton wieder *loslässt*, und davor, schon beim Herunterdrücken, wird das JavaScript-Event `onmousedown` ausgelöst. Für dieses Event ist nun im HTML-Code des Links eine spezielle Routine hinterlegt worden, welche eine JavaScript-Funktion namens `rwt` aufruft, der verschiedene Parameter übergeben werden – das ist die durch Kommata separierte Liste in der langen runden Klammer direkt hinter `rwt`.

Das Kürzel `rwt`, das wenig verhohlen für »rewrite« steht, bezeichnet eine von Google weiter oben im Quelltext bereitgestellte JavaScript-Funktion – das ist eine kleine Programmroutine, die von der Website im Browser hinterlegt wird, um bei Bedarf (hier beim Anklicken eines Links) *nach* dem eigentlichen Ladevorgang der Seite ausgeführt zu werden. Sie ist so programmiert, dass sie im Moment des Klicks den `href`-Parameter des Links, also die Ziel-URL, überschreibt und in die folgende längliche Adresse ändert:

```
https://www.google.de/url?sa=t &ct=j &q= &esrc=s &source=web &cd=1
&cad=rja &uact=8
&ved=0ahUKEwj5t3t0YLZAhULLSwKHb-eAYcQFgggMAA
&url=http://www.tagesschau.de/
&usg=A0vVaw1_8NtrWpEbLefVcTd78eDR18
```

Zwischen dem Drücken des Mausbuttons und dem Loslassen *wird also das Ziel des Links ausgetauscht*. Man ruft deshalb tatsächlich eine Google-URL auf (nämlich `www.google.de/url?...`), und erst von dort aus wird man mit einem HTTP- (Hypertext Transfer Protocol-) Weiterleitungsmechanismus (HTTP status code 302) auf das eigentliche Ziel, also `www.tagesschau.de`, weitergeleitet.⁹ Diese Operation erfolgt bei einer normalen Internetverbindung in Millisekundschnelle, so dass die Nutzer_in nicht bemerkt, dass sie, bevor sie am Ziel ankommt, noch kurz `www.google.de/url?...` besucht.¹⁰

8 | Zur besseren Lesbarkeit wurden Leerzeichen eingefügt und der %-encodierte URL-Parameter `url=http%3A%2F%2Fwww.tagesschau.de%2F` ersetzt durch die decodierte Form `url=http://www.tagesschau.de/`.

9 | Wenn der von Google selbst entwickelte Chrome-Browser verwendet wird, kommt dieser *rewrite*-Mechanismus nicht zum Tragen, sondern die Hyperlinks auf der SERP enthalten ein HTML-Attribut `ping="/url?..."`, welches beim Anklicken des Links zeitgleich zu dem Aufruf der Zieladresse die `google.de/url?...` Seite mit aufruft, ohne dafür extra den `href`-Parameter des Links austauschen zu müssen. Das `ping`-Attribut wurde erst mit HTML-Version 5 eingeführt und dient dazu, eine »transparentere« Infrastruktur für das Tracken von Klicks auf externe Links zu schaffen. Auch Firefox unterstützt das `ping`-Attribut, Google verwendet jedoch bei diesem Browser den *rewrite*-Mechanismus, weil Firefox-User den `ping`-Mechanismus per Konfiguration ausschalten können. Vgl. Oxley 20.03.2014.

10 | Jede_r kann das selbst überprüfen durch folgende Tricks: Variante 1: SERP aufrufen, Internetverbindung kappen, dann auf den Link klicken. Variante 2: Mit der Maus auf den Link klicken, den Mausbutton aber nicht loslassen, sondern die Maus mit gedrückter Taste vom Link

Bei diesem Kurzbesuch auf einem Google-Server werden nun allerhand Informationen an Google übertragen. Das ist der oben abgedruckten URL zu entnehmen: Hinter dem ? und hinter jedem &-Zeichen in der URL beginnt der Name einer Variable (eines »Parameters«). Diesen Variablen wird hinter dem jeweiligen =-Zeichen jeweils ein bestimmter Wert zugewiesen und mit dem Aufruf der URL werden diese Werte dann an den Server übertragen.¹¹ Das heißt, beim Anklicken eines Suchresultats übermittelt die Resultateseite an Google bestimmte Informationen zurück, in diesem Fall handelt es sich um elf verschiedene Parameter. Darunter – am offensichtlichsten – ist der Parameter `url=http://www.tagesschau.de/`. Er übermittelt an Google, *welches* Suchresultat angeklickt wird. Der Parameter `cd=1` zeigt an, an welcher Stelle sich das angeklickte Suchresultat auf der Resultateseite befunden hat. In diesem Fall war es das oberste Resultat, daher der Wert 1, doch der Wert von `cd` erhöht sich für jedes weiter unten gelistete Resultat auf der SERP um 1. Google ist somit nicht nur in der Lage, Statistiken darüber zu führen, welche Resultate angeklickt werden, sondern auch darüber, wie das Anklickverhalten davon abhängt, an welcher Stelle ein Resultat auf der SERP gelistet wird.

Im Allgemeinen ist es nicht bekannt und gilt als gut gehütetes Betriebsgeheimnis, welche Informationen die verschiedenen übermittelten Parameter genau codieren. Einzelnes darüber lässt sich jedoch auf dem Wege des *reverse engineering* und durch technische Einblicke herausfinden oder erraten.¹² So führt der Parameter `usg` etwa eine verschlüsselte Version des `url`-Parameters und erzeugt somit lediglich eine Informationsredundanz, die dazu dienen kann, Verfälschungen in der Übermittlung zu erkennen. Interessanter ist dagegen diese Variable:

```
ved=0ahUKEwjJ5t3t0YLZAhULLSwKHb-eAYcQFggUAA
```

Ihr Wert ist ein 40 Zeichen langer String (Zeichenkette), der aus mehreren Bestandteilen zusammengesetzt ist: Die ersten 7 Zeichen, die mittleren 25 Zeichen und die hinteren 8 Zeichen sind für sich jeweils ein Teilstring. Mit Probieren und *educated guessing* findet man heraus, dass der hintere Teil detaillierte Informationen darüber encodiert, wo auf der Resultateseite der angeklickte Link positioniert war. Insbesondere werden hier qualitative Informationen erfasst, zum Beispiel

wegziehen, danach erst loslassen. Durch diese Prozedur wird der Link nicht aufgerufen (das wäre nur beim Loslassen der Taste *auf* dem Link der Fall), aber dasonmousedown-Event wird trotzdem ausgelöst, so dass die Zieladresse des Links ausgetauscht wird. Führt man erneut mit der Maus über den Link, zeigt die Statuszeile des Browsers das modifizierte Ziel.

11 | Es handelt sich bei den Werten dieser URL-Parameter um die Daten, die zuvor der Funktion `rwt` als Argumente übergeben und somit von Google selbst für jede individuelle Suchanfrage in den Quelltext der SERP hineingeschrieben werden.

12 | Siehe für Details: Ny 08.06.2016; Resnik 22.05.2013; Wittersheim 31.03.2016; sshay77 18.07.2015.

ob es sich um einen *Ad link*, um ein *image result*, um einen Eintrag im *knowledge graph* oder um ein herkömmliches *organic search result* handelte.¹³

Was die ersten 7 Zeichen des *ved*-Parameters encodieren, entzieht sich der Kenntnis des Verfassers. Im Augenmerk soll nun jedoch der 25 Zeichen lange mittlere Teil stehen (oben fett gedruckt). Wie sich zeigt, ist dieser Mittelteil ein Identifikationscode, der es erlaubt, die *einzelne Suchsession* zu identifizieren: Ermittelt man etwa die *ved*-Parameter der ersten drei Suchresultate auf der Resultatseite unserer Suche nach »aktuelle Nachrichten«, dann sehen sie so aus:

1. Resultat: ved=0ahUKEwj**j5t3t0YLZAhULlSwKHb-eAYc**QFgguMAA
2. Resultat: ved=0ahUKEwj**j5t3t0YLZAhULlSwKHb-eAYc**QFgg7MAE
3. Resultat: ved=0ahUKEwj**j5t3t0YLZAhULlSwKHb-eAYc**QFghHMAI

Der Mittelteil bleibt für alle Links identisch; der hintere Teil verändert sich (denn er gibt ja die Position des Links auf der Seite an). Startet man nun eine neue Suche nach »aktuelle Nachrichten« oder nach einem anderen Stichwort, indem man den Browser schließt, dann wieder öffnet und *google.de* neu aufruft, dann erhält man folgende *ved*-Parameter der ersten drei Suchresultate:

1. Resultat: ved=0ahUKEw**ikooTQ6oTZAhWQ_aQKHb**eFBAsQFgguMAA
2. Resultat: ved=0ahUKEw**ikooTQ6oTZAhWQ_aQKHb**eFBAsQFgg7MAE
3. Resultat: ved=0ahUKEw**ikooTQ6oTZAhWQ_aQKHb**eFBAsQFghHMAI

Es fällt auf: Der Mittelteil ändert sich mit der neuen Suchanfrage, sogar wenn nach den gleichen Stichworten gesucht wird. Das deutet darauf hin – und mit weiteren Tests lässt sich dies erhärten –, dass der Mittelteil von *ved* dazu genutzt werden kann, die konkrete *Suchsession* eindeutig zu identifizieren. Jede einzelne Suchanfrage wird auf diese Weise mit einer eindeutigen Kennung versehen.¹⁴ Immer wenn eines der Suchresultate angeklickt wird, registriert Google also nicht nur, welche Ziel-URL angeklickt wurde und wo auf der SERP diese aufgeführt war, sondern das angeklickte Resultat kann der einzelnen Suchanfrage wieder zugeordnet werden. Dadurch lassen sich ganze Suchverläufe serverseitig erfassen.¹⁵ Es

13 | Es handelt sich bei diesem hinteren Teil von *ved* um den »alten« *ved*-Parameter. Es scheint in den letzten Monaten eine Veränderung gegeben zu haben: Während ursprünglich die Session-Identifikation durch einen separaten Parameter ei möglich war und *ved* dann nur die Angaben zur Position des angeklickten Links auf der Seite speicherte, ist ei nun entfallen, aber der jetzt deutlich längere *ved*-Parameter kann seine Funktion übernehmen. Auf die Informationen, die in diesem hinteren Teil encodiert werden, kann hier nicht näher eingegangen werden, siehe ausführlich Resnik 22.05.2013; Kelly 02.01.2014.

14 | Reproduziert man diese Experimente, weichen die konkreten Werte des Mittelteils natürlich von den hier dargestellten ab. Sie werden für jede Suchanfrage neu vom Server vergeben.

15 | Auch wenn, anstatt ein Resultat anzuklicken, die Suche verfeinert wird (indem ein neues oder zusätzliches Suchstichwort eingegeben und wieder auf »Suchen« geklickt wird), wird ein Session-Informationscode an den Server übertragen. Dies erlaubt es, die Verknüpfung der ersten

kann zum Beispiel ausgewertet werden, wie eine mehrschrittige Suche verläuft, auf welches Resultat in welchem Schritt dieses Suchverlaufs geklickt wird, nach welchen Stichworten in Reaktion auf vorherige Resultate gesucht wird und vieles mehr.

Mit diesen Mechanismen lässt sich also detailliert das Suchverhalten von Nutzer_innen erfassen. Dabei geht es einerseits um eine Bemessung der »Relevanz« der Suchresultate: Welche Stichworte werden gesucht, wie ist der Verlauf der Anfragen und welchen Resultaten wird gefolgt. Andererseits geht es aber auch um die metrische Analyse der Nutzer-Responsivität auf das Design der *Aufbereitung* der Resultate auf der Resultatenseite. Es wird erfasst, ob eher *organic search results* oder andere Seitenelemente (Ads, *knowledge graph*, *image search*) bevorzugt werden. So existieren etwa detaillierte Auswertungen, welche Positionen auf der Google-SERP statistisch am ehesten angeklickt werden.¹⁶

Durch diese technischen Mechanismen, die im Hintergrund der Resultatenseite operieren, werden eine große Menge Nutzungsdaten unfreiwillig und unbemerkt erhoben. Auch wenn diese nicht grundsätzlich in die Kategorie »personenbezogene Daten« fallen, besitzen sie einen enormen wirtschaftlichen Wert. Doch damit nicht genug – Google ist nämlich tatsächlich auch an personenbezogenen Daten interessiert. Ist man im selben Browser, etwa in einem anderen Tab oder in einem anderen Fenster, zeitgleich zu einer Google-Suchanfrage in einem anderen Google-Service eingeloggt – etwa in Gmail, in Google Drive oder einem der zahlreichen weiteren Services des Unternehmens –, dann wird von Google im Browser ein Cookie hinterlegt, welches die Nutzer_in anhand ihres »Google Accounts« eindeutig identifiziert und auch bei der Google-Suchanfrage mit an den Server übertragen wird.¹⁷ Die Erhebung des detaillierten Suchverlaufs und des Klick-Verhaltens kann in diesem Fall also nicht nur einem *anonymen User*, sondern einem bekannten Nutzeraccount zugeordnet werden.

Das ist ein großer qualitativer Schritt: Während der anonyme User (im Google-Jargon »*client*« genannt) nur über eine Nutzungssession hinweg verfolgt werden kann, erlaubt die Zuordnung zu einem Nutzeraccount (im Google Jargon: *user*) erstens, die anfallenden Daten über das Suchverhalten mit den Daten zu verknüpfen, die im Rahmen aller anderen Google-Services gespeichert werden –

Suchanfrage mit der zweiten vorzunehmen. Die Übertragung dieses Parameters bei einer neuen oder verfeinerten Suche wird über ein verstecktes (hidden) Formularfeld gelöst, welches einen Identifikationsparameter namens *ei* überträgt:

```
<input value="JXBzWpyYLI0SsAeup7cY" name="ei" type="hidden">
```

16 | Siehe etwa Mediative 2014.

17 | Auch wenn die Nutzer_in in dem Browser einmal eingeloggt war und sich dann ausgeloggt hat, bleiben Cookies bestehen, die sie eindeutig identifizieren können. Cookies müssten nach jeder Sitzung vollständig gelöscht werden, um diesen Effekt zu umgehen.

darunter E-Mail-Inhalte und Dokumente in Google Drive, Standort des Android-Smartphones, Adressbuch, Telefonanrufe, SMS-Nachrichten. Es erlaubt zweitens, verschiedene Nutzungssessions, die nacheinander oder auf verschiedenen Geräten erfolgen, miteinander zu verknüpfen und auf diese Weise lebenslange und geräteübergreifende Suchhistorien anzulegen. In diesem deutlich höherdimensionalen Datenraum kann zum Beispiel die individuelle Responsivität für verschiedene Suchresultate mit den aktuellen E-Mail-Inhalten, mit dem Standort des Android-Telefons, mit dem YouTube-Video, das gerade angeschaut wurde, mit den Daten, die das Fitnessarmband über einen Health-Service erfasst hat, korreliert werden. Solche Verknüpfungen und Korrelationsanalysen (Data-Mining) können auch nachträglich, nach einigen Jahren oder durch andere Unternehmen, an die die Rohdaten verkauft werden, vorgenommen werden. Es ist möglich, aus diesen Daten detaillierte psychologische und affektologische Profile von Nutzer_innen anzufertigen und als abgeleitete Daten in Form von *eScores* für Zwecke des Risikomanagements (zum Beispiel zur individuellen Bepreisung von Krankenversicherungen, bei Einstellungsverfahren auf dem Job-Markt, zur Quantifizierung von Kreditwürdigkeit, Bonität oder Rückfallwahrscheinlichkeiten bei Kriminalprozessen) oder der individualisierten Werbung zu vermarkten (O'Neil 2016).¹⁸

Beispiel 2: Single-Sign-on-Services

Eine zweite Klasse von Beispielen für die Erfassung scheinbar freiwillig bereitgestellter Personen- und Nutzungsdaten auf vernetzten Plattformen betrifft Techniken der Ansprache, der Gestaltung von Dialogboxen und Interfaces, die man als »Nudging« bezeichnen kann. Im Gegensatz zu dem vorangegangenen Beispiel bestehen diese Fälle nicht aus technischen Einrichtungen, die verdeckt operieren, sondern es handelt sich hier nun um Konstellationen, in denen an der Oberfläche arbeitende Design- und Gestaltungsstrukturen mit einem bestimmten Nutzungsverhalten oder einer bestimmten *Disposition*, sich in seinem Nutzungsverhalten beeinflussen zu lassen, Hand in Hand arbeiten.

18 | Häufig wird bei diesem Thema auf den Umstand hingewiesen, dass die Suchresultate, die zum Beispiel bei der Suche »aktuelle Nachrichten« angezeigt werden, sich plötzlich verändern, wenn man sich im Hintergrund in seinen Google-Account einloggt. Denn Google verwendet dann auch die aus E-Mail-Inhalten oder anderen Services über die Nutzer_in bekannten Daten, um für diese Nutzer_in individuell möglichst »relevante« Resultate und Werbeanzeigen anzuzeigen. Ich konzentriere mich in diesem Artikel gezielt auf die weniger bekannte und weniger beachtete andere Richtung dieser Feedbackschleife: Nicht nur was man angezeigt bekommt verändert sich abhängig von den Datenspuren, die man hinterlassen hat, sondern die eigenen Klicks hinterlassen Datenspuren, die die Anfertigung detailreicher psychologischer und behavioreller Metriken erlauben.

Ein verbreitetes Beispiel für das, was in diese Klasse fällt, sind die Authentifizierungsdienste »Google Sign-In« und »Facebook Login«. Das sind von den beiden Konkurrenten Google und Facebook angebotene Services, die sich in beliebige Android- und iPhone-Apps sowie auf Websites integrieren lassen und mithilfe derer die Entwickler_innen einer App oder Website es ihren Usern ermöglichen, sich mit ihrem Google- bzw. Facebook-Account bei der App oder Website zu registrieren, anstatt mit einem selbst gewählten Benutzernamen und Passwort für diesen Service einen neuen User-Account anzulegen. Googles und Facebooks Authentifizierungsschnittstellen sind extrem verbreitet, man findet die Möglichkeit eines »Login with Facebook« oder eines »Sign in with Google« auf Shopping-Plattformen, bei Dropbox und Doodle, in Dating-Portalen wie Tinder oder OkCupid, auf Nachrichtenportalen wie Spiegel Online oder bild.de, bei Airbnb, Uber, Netflix, Spotify und SoundCloud, um nur einige sehr populäre Services zu nennen.

Die grundsätzliche Idee von Single-Sign-on-Schnittstellen ist, dass es aus User-Sicht als mühevoll gilt, durch Angabe einer E-Mail-Adresse und eines Passworts für jeden Service einen eigenen Benutzeraccount anzulegen. Sowohl bei der initialen Registrierung (»onboarding«¹⁹) für einen neuen Benutzeraccount auf einer Website oder in einer App, wo im herkömmlichen Verfahren oft noch eine zusätzliche Schleife zur Überprüfung der E-Mail-Adresse oder Telefonnummer gefahren werden muss, als auch bei der täglichen Benutzung gilt der herkömmliche Login als ein möglicher Reibungspunkt, zum Beispiel weil User ihre Passwörter leicht wieder vergessen können, besonders wenn sie für viele verschiedene Seiten jeweils verschiedene Zugangsdaten verwenden.

Die Möglichkeit eines zentralen, auf einen Klick erfolgenden Logins via Google oder Facebook ist aber nicht nur bequemer, sondern bietet auch für die Betreiber_innen des Services Vorteile. Neben dem »frictionless onboarding« besteht ihr Hauptgewinn daraus, dass im Moment eines Sign-ins via Google oder Facebook zahlreiche personenbezogene Daten über den User an den Service übertragen werden, die auf dem Wege einer herkömmlichen Registrierung mit sehr viel mehr Aufwand den Nutzer_innen entlockt werden müssten. So überträgt Google Sign-In standardmäßig mindestens den vollen Klarnamen, die (verifizierte) E-Mail-Adresse und ein Bild des Benutzers. Zugriff auf weitere über den Nutzer hinterlegte Informationen, zum Beispiel im Google+-Profil oder über YouTube, ist prinzipiell möglich. Bei Facebook erhält die Website, auf der man sich mittels Facebook Login anmeldet, standardmäßig Zugriff auf alle allgemein zugänglichen Facebook-Profildaten (darunter Name, Foto, E-Mail-Adresse, Altersklasse, Gender, Locale, Zeitzone), sowie die Liste der Facebook-Freunde des Users, die ebenfalls diesen Service nutzen. Facebook ermöglicht es den Entwicklern einer

19 | »Onboarding« ist im UX-Jargon der Prozess der Gewinnung und Registrierung eines neuen Users für einen Service. Es entscheiden oft wenige Unannehmlichkeiten in der Benutzerführung darüber, ob sich die Nutzer_in vor Abschluss einer vollständigen Registrierung noch abwendet.

Abbildung 1: Login-Dialog auf airbnb.com. Quelle: Screenshot des Verfassers vom 30.01.2018.

Website oder App prinzipiell, noch auf viele weitere Nutzerdaten zuzugreifen (darunter beispielsweise die Like-Liste, die Liste der Freunde, Geburtsdatum, Aufenthaltsort, Beziehungsstatus, Arbeitsleben und Berufsqualifikationen, Freizeitinteressen etc.²⁰). Facebook verlangt allerdings bei Zugriff auf solche erweiterten Datensätze, dass dafür von der Nutzer_in eine einmalige explizite Zustimmung eingeholt wird. Um eine Metapher zu bilden: Der Verwendung des Facebook Logins würde in der »realen Welt« entsprechen, wenn man stets beim Betreten eines Shops die Liste seiner Interessen, Freunde, Likes, Statusposts, Sprachfähigkeiten, Schulabschlüsse, Berufserfahrungen etc. am Eingang abgeben würde. Der Shop könnte dann ganz schnell die räumliche Anordnung seiner Produktregale, die Anordnung der Produkte in diesen Regalen, sowie die Preise der Produkte und mögliche Sonderangebote auf diesen einen Benutzer abstimmen. Das gleiche gilt auch, wenn man ein Versicherungsbüro betritt oder sich um einen Job bewirbt – um nur wenige Felder zu nennen, in denen diese Daten verwertet werden.

Natürlich könnte man hier einwenden, dass die User sich wissentlich für die Verwendung des Sign-in-Services entscheiden. Die meisten (wenn auch nicht alle) Websites und Apps, die einen Sign-in via Facebook oder Google anbieten, stellen die Option zur Verwendung dieses zentralen Logins in der konkreten Dialogbox, die zur Registrierung oder Anmeldung auffordert, neben der Option einer her-

kömmlichen Registrierung mit E-Mail-Adresse und Passwort dar (siehe 1). Außerdem verpflichten sowohl Google als auch Facebook die Entwickler_innen externer Seiten und Apps dazu, die mit dem Sign-In eingeholten Berechtigungen zu einem Zugriff auf Nutzerdaten explizit zu nennen und explizit Zustimmung dafür einzuholen. Es liegt hier also vermeintlich alles transparent zu Tage und ist der ›freien‹ Entscheidung der Nutzer_innen überlassen. Und doch, oder gerade deshalb, können die zentralen Authentifizierungsservices in mehreren Hinsichten als ein subtiler Nudge bezeichnet werden – als eine Technik der zwanglosen Beeinflussung von Nutzerentscheidungen durch das Design von »Wahlarchitekturen« (*»choice architectures«*²¹) auf der Grundlage verhaltenswissenschaftlicher und psychologischer Erkenntnisse.²² Das verrät auch ein Blick in die *Developer Guidelines* von Facebook zum Facebook Login.²³ Dort heißt es im Abschnitt »User Experience Design«:

»The onboarding experience is one of the most important user experiences in your app. A high quality onboarding experience can lead to conversion rates²⁴ above 90% and encourages people to become more engaged and profitable.« (Ebd.)

Es werden von dieser Annahme ausgehend detaillierte Tipps für die beste Gestaltung einer Login-Seite gegeben, die eine möglichst hohe »Konversionsrate« erzielen können. Diese Hinweise erwecken keineswegs den Anschein, dass Facebook sich den durchschnittlichen User als ein rational, voll informiert und bewusst entscheidendes Individuum vorstellt:

»Reducing unnecessary steps is one of the most effective ways to improve your conversion rate. Avoid asking users to first tap ›Login‹ or ›Register‹ to get to the Facebook login button. With Facebook Login, this is an unnecessary step. There's no need for people to even have stop to think about if they have an account or not.

20 | Siehe die vollständige Liste verfügbarer Datenfelder und ihrer *permission scopes*: <http://developers.facebook.com/docs/facebook-login/permissions/>; Stand: 2018-02-10.

21 | Vgl. zur ursprünglichen Verwendung dieses Begriffs in der Verhaltensökonomie: Thaler und Sunstein 2008. Zur Übertragung auf digitale Benutzerschnittstellen als »Interface Nudges« siehe Mühlhoff 2018.

22 | Avi Charkham (25.08.2012) analysiert in einem Blog-Post mit Screenshots die Design-Tricks, die Facebook bei der Gestaltung der Benutzeroberfläche zum Erteilen von Zugriffsrechten für externe Apps verwendet.

23 | <https://developers.facebook.com/docs/facebook-login/userexperience>; Stand: 2018-02-10.

24 | Mit »conversion rate« ist die relative Anzahl neuer App-Benutzer oder Seitenbesucher gemeint, die auch tatsächlich einen Nutzeraccount anlegen, bei denen das Onboarding also erfolgreich verläuft.

In addition, after people have logged in with Facebook, don't prompt them to create a username or password. One of the most popular reasons people log in with Facebook is because «it's fast and easy and I don't have to enter a password». After logging in with Facebook, people especially do not want to have to create a username or password.» (Ebd.)

Auch der Sensibilität persönlicher Informationen sind die Autor_innen dieser Guidelines sich bewusst, zumindest insofern als ein zu »forsches« Vorgehen beim Einholen von Berechtigungen sich negativ auf die »conversion rate« auswirke. Ein probates Mittel, wieder aus dem Repertoire des Nudging, wird wenige Absätze später an die Hand gegeben:

»Only ask for the permissions you need

The fewer permissions you ask for, the easier it is for people to feel comfortable granting them. We've seen that asking for fewer permissions typically results in greater conversion. You can always ask for additional permissions later after people have had a chance to try out your app. [...] People are most likely to accept permission requests when they understand why your app needs that information to offer a better experience. So trigger permission requests when people are trying to accomplish an action in your app which requires that specific permission.» (Ebd.)

Anders gesagt heißt das, man solle die von der App oder der Website geforderten (über das Default hinausgehenden) Zugriffsrechte auf persönliche Informationen nicht zu Beginn, bei der Registrierung mittels Facebook Login, offenlegen und die Zustimmung bereits dann abfragen, sondern jede benötigte zusätzliche Berechtigung solle besser »in context« eingeholt werden – also nachdem sich die Nutzer_in bereits darauf eingelassen hat, im Rahmen des Onboarding einen ersten Teil ihrer Informationen preiszugeben und vielleicht schon verschiedene Eingaben oder Bewegungen auf der Seite beziehungsweise in der App vollzogen hat, wonach die Wahrscheinlichkeit geringer ist, dass sie alles wieder abbricht. Wird eine Berechtigung, die man eigentlich nicht so gerne erteilt, erst später eingeholt – zum Beispiel nachdem der User aufwendig ein Profil konfiguriert und Texte über sich hinterlegt hat (etwa auf einer Dating-App oder beim Einreichen eines Inserats auf einer Verkaufsplattform) oder kurz vor Ende eines Bezahlvorgangs –, dann ist die Wahrscheinlichkeit, dass die Berechtigung nicht erteilt wird, geringer als wenn man den User direkt zu Beginn danach fragt. Genau zu diesem Stil eines »User Experience Designs« zwischen künstlich-oberflächlicher Wohlfühl-atmosphäre und subtiler Entmündigung fordert Facebook die Entwickler_innen externer Apps und Seiten in den Guidelines aktiv auf. Das Design von Benutzeroberflächen im Netz kreist heute zu großen Teilen darum, statistisch gesehen möglichst effiziente »Wahlarchitekturen« zur Produktion hoher Konversionsra-

ten bereitzustellen.²⁵ UX Design, das zeigt sich hier gut, operiert stets unter bestimmten psychologischen Prämissen, die aber nicht rein deskriptiv verwendet werden, sondern – als langfristiges Resultat vieler kleiner solcher Ansprachen durch entmündigende Interfaces – auch eine bestimmte Subjektivität, eine bestimmte Haltung des »Ein Passwort eingeben ist mir zu lästig«, »Ich will überhaupt nicht darüber nachdenken müssen, ob ich auf dieser Seite schon einen Account habe oder nicht« *produzieren*.

Nun mag bei alledem die Frage auftauchen, warum Facebook und Google eigentlich motiviert sind, das Leben externer App- oder Website-Entwickler_innen durch die Bereitstellung solcher aufwendigen Programmierschnittstellen und dazugehöriger UX-Guidelines und durch die Weitergabe von Kundendaten an Dritte zu erleichtern. Hier zeigt sich nun eine Art »doppeltes Nudging« im Zusammenhang mit den zentralen Authentifizierungsdiensten, denn es nudgen nicht nur die Website-Betreiber ihre User zu einer Entscheidung, durch die sie an deren persönliche Daten kommen, sondern *es nudgen auch Facebook und Google die Entwickler_innen von Apps und Websites* dazu, diese Authentifizierungsservices zu benutzen. Der Punkt ist, dass Google und Facebook von der Bereitstellung der Authentifizierungsdienste insofern profitieren, als sie damit an weitere enorm detaillierte Informationen über das Konsumverhalten ihrer Nutzer_innen gelangen. In dem Moment, in dem ein User sich über den Facebook Login bei einem externen Service (beispielsweise eine Dating-App) registriert, erhält nicht nur dieser Service die personenbezogenen Daten der Nutzer_in, sondern auch Facebook die Echtzeit-Information, dass diese Nutzer_in gerade diesen externen Service nutzt. Hier kommt nun der Plattformeffekt zum Tragen, der bedeutet, dass dies für Facebook und Google umso interessanter wird, je verbreiteter ihr Login-Verfahren bei externen Services ist, denn diese Plattformen verfügen dann – anders als die einzelnen Anbieter externer Services – über das Wissen, welche *Kombination* verschiedener externer Services ein bestimmter Benutzer gleichzeitig oder nacheinander verwendet. Gegen den für die Plattformen sehr geringen Preis der Überlassung von persönlichen Daten über die User erhält die Plattform also detaillierte Informationen darüber, was der User gerade *außerhalb* der eigenen Plattformgrenzen tut – ob er einen Mietwagen gebucht hat, wie häufig er die Online-Dating-App öffnet, ob er bei Airbnb etwas angeboten oder gesucht hat etc.

25 | Unten auf der zitierten Seite bietet Facebook den Entwickler_innen auch sein eigenes Statistik-Tool zum Erfassen der Konversionsrate an: »Facebook Analytics lets you monitor your conversion rates for free« (ebd.) und ermöglicht die detaillierte Auswertung, an welchem Schritt im Benutzungsfluss die Nutzer_innen verloren gehen. Das heißt insbesondere auch, dass Websites, welche den Facebook Login verwenden, die schrittgenaue Erfassung von Nutzerbewegungen durch Facebook ermöglichen. Diese Daten werden somit durch Facebook erfasst und können auch mit anderen Profildaten und daraus abgeleiteten Persönlichkeitsprofilen und Metriken kombiniert werden. Siehe auch Mühlhoff 2018.

Echtzeit-Informationen über das, was Nutzer_innen gerade tun, gelten als extrem wertvolle und gut vermarktbarere personenbezogene Daten.

Aus Sicht der Plattformunternehmen Google und Facebook sind die zentralen Authentifizierungsdienste also eine Technologie, mit der die Reichweite der eigenen Tracking- und Datensammelmechanismen über die Grenzen der eigenen Plattform hinaus enorm erhöht wird. Diese höhere Penetration ihrer dezentralen Trackingdienste im Netz macht diese Dienste profitabel, denn die erhöhte Reichweite erlaubt die Anfertigung detaillierterer behavioreller und psychometrischer Analysen der Nutzer_innen durch Big Data-Verfahren, die in einem breiten Kontext von Anwendungsmöglichkeiten, vom Risk Controlling über *health scores* bis zu *targeted ads*, vermarktet werden können. Gerade für Facebook, das anders als Google zunächst eine geschlossene Plattform ist, ist diese Vergrößerung der Reichweite über die eigenen Plattformgrenzen hinaus von enormer Bedeutung.²⁶

An dieser zweiten Klasse von Beispielen ist eine Technik der ›freiwilligen‹ und ›konsensuellen‹ Gewinnung von Nutzerdaten erkennbar, die nicht wie in der ersten Klasse unbemerkt und unter der Oberfläche verfährt, sondern auf Grundlage expliziter situativer Nutzerentscheidungen möglich ist und über das Design von *choice architectures* statistisch gesehen möglich gemacht wird. Anders als bei einer grundsätzlichen Einwilligung zum Beispiel in »Allgemeine Nutzungsbedingungen« ist hier eine fallweise, aktive Handlung der Nutzer_in erforderlich. Doch es kann kaum davon die Rede sein, dass die Nutzer_in sich grundsätzlich in einer wissenden oder urteilskompetenten Position befindet, denn es bleibt strategisch verborgen, welche Formen der Datenerhebung und Datenaggregation und welche Generierung abgeleiteter Daten diese situative Zustimmung ermöglicht.

Dieses subtile Framing, das ein freiwillig unfreiwilliges, konsensuelles, aber doch nur oberflächlich informiertes Nutzerverhalten erzeugen möchte, ist die bedeutendste Technik im Zusammenhang mit der massenweisen Erhebung von Nutzerdaten im Netz. User Tracking ist dabei nicht nur das Geschäft großen Firmen und einer weltweiten Kohorte von Techniker_innen und Ingenieur_innen, sondern auch der Gegenstand eines Diskurses in Verhaltenswissenschaften und angewandter (Verhaltens-)Psychologie. Benutzerführung und »UX Designs« werden speziell auf die *Trackbarkeit* der User optimiert, das heißt, die Gestaltung von Benutzeroberflächen wird darauf ausgerichtet, dass Nutzer_innen möglichst viele verwertbare Daten hinterlassen. Das Wissen, die Techniken, die Infrastrukturen, Narrative und Subjektivitäten, die im Zusammenhang mit diesen Praktiken entstehen, bilden im vollen Sinn ein *Dispositiv* (vgl. Foucault 1978 [1977]). Das heißt

26 | Facebook Login ist nicht die einzige Technologie, die das erlaubt. Noch wichtiger und als fundamentaler *game changer* in der Internetwelt betrachtet war die Einführung des »Like-Buttons« auf externen Seiten, die es Facebook ebenfalls erlaubt, die Aufrufe externer Seiten zu loggen (ohne dass der Like-Button dafür tatsächlich betätigt werden müsste, weil schon das Laden des Code Snippets, der ihn zur Anzeige bringt, einen Kontakt zum Facebook-Server aufbaut und somit den Seitenaufruf verrät).

insbesondere zweierlei: Erstens kann man ihre gesellschaftlichen Auswirkungen oder Gefahren nicht auf der Ebene von Einzelfällen beurteilen, sondern muss das implizite strategische Zusammenspiel vieler Orte, an denen solche Interaktionsdesigns verwendet werden, und der in der Schaffung dieser Interfaces beteiligten Wissenspraktiken und ökonomischen Interessen in den Blick nehmen. Zweitens spielt in der Funktionsweise dieses Dispositivs die Hervorbringung einer bestimmten subjektiven Wahrnehmungsweise der Nutzer_innen selbst eine entscheidende Rolle, denn die Datenerhebung erfolgt hier nicht heimlich und unter der Oberfläche, sondern durch die aktive und prinzipiell wissende Mitwirkung jedes Einzelnen.

Beispiel 3: Menschengestützte Künstliche Intelligenz

In der ersten Fallstudie wurde diskutiert, wie Google die Reaktionen seiner Suchmaschinenutzer_innen auf angezeigte Suchresultate erfasst und mit zahlreichen weiteren über die konkrete Nutzer_in bekannten persönlichen, psychologischen und wirtschaftlichen Daten korrelieren kann. Anhand dieser Daten kann Google die Qualität seiner Suchresultate, seiner *risk-assessment*-Services und des Targetings seiner individualisierten Werbung verbessern. Meine These ist, dass hier eine technologische Strategie sichtbar wird, die für gegenwärtige Netztechnologien paradigmatisch ist. Etwas allgemeiner kann dieses Paradigma so formuliert werden: *Menschliche kognitive Kapazitäten gelten als Ressourcen, die unbemerkt in einen technischen Apparat eingespannt werden können, der im Ganzen dadurch eine bestimmte informationsverarbeitende Aufgabe optimieren kann.* Es handelt sich hierbei um einen allgemeineren technischen Trend, den ich im Folgenden als »menschengestützte künstlichen Intelligenz«, oder *Human-Aided AI*, bezeichne (Mühlhoff 2019).²⁷

Im Hintergrund dieser Überlegung steht die Beobachtung, dass sich mit dem Web 2.0 eine schleichende, aber grundlegende Transformation im Verständnis von künstlicher Intelligenz (KI) zugetragen hat. In der Mitte des 20. Jahrhunderts, zu Zeiten von Turing und Minsky, verstand man unter KI die Technikvision, dass eine Rechenmaschine irgendwann zu kognitiven Leistungen fähig sein würde, die die kognitiven Leistungen eines Menschen *ersetzen* können (Simulationsverständnis von KI). Diese Überlegung führte etwa zu Konstrukten wie dem Turing Test oder des Chinese Room-(Gedanken-)Experiments. Im beginnenden 21. Jahrhundert, das weitgehend vom Gedanken der Vernetzung getragen ist – nicht nur des Sozialen, sondern auch in Gestalt dezentraler Rechenkapazitäten

27 | Im Englischen ist der Term »human-assisted artificial intelligence« geläufiger. Obwohl er noch nicht in der wissenschaftlichen Literatur angekommen ist oder zum Gegenstand kritischer Debatten wurde, wird er in der Blog-Sphäre rege verwendet, siehe exemplarisch Pichsenmeister 02.12.2016.

und Informationsflüsse –, hat sich das Verständnis von KI diversifiziert und weiterentwickelt. Künstliche Intelligenz bezeichnet heute nicht mehr nur Routinen oder Softwareprogramme, die die kognitive Leistung des Menschen simulieren können, sondern hat auch die Gestalt eines dezentralen technischen, wirtschaftlichen, sozialen und politischen Apparats angenommen, der die kognitiven Mikrofähigkeiten von Menschen möglichst passgenau *einhegt* und *abschöpft*, um im Ganzen – als eine auf höherer Ebene emergierende Form der KI – eine bestimmte informationsverarbeitende Leistung zu erbringen. Das simulatorische Verständnis von KI, so die kritische These, wird in der Praxis heute in vielen Bereichen durch ein *Immersionsverständnis von KI* ersetzt oder ergänzt. Menschengestützte KI zeichnet sich dadurch aus, dass menschliche kognitive, soziale und affektive Ressourcen lückenlos in ein größeres Gefüge von Computernetzwerken eingebaut werden, als »wet-ware« in einem heterogenen Ensemble von Hardware und Software, Menschen und Maschinen, etwa um Trainingsdaten zu gewinnen, die Fehleranfälligkeit von KI zu reduzieren oder um die menschliche kognitive Fähigkeit als fest verschaltete Ressource in hybriden Mensch-Maschine-Netzen auszuheben.

Im Zusammenhang mit Internetanwendungen finden sich zahlreiche große und kleine Beispiele für diese Technik einer *Human-Aided AI*. Sweatshops auf den Philippinen, in denen schlecht bezahlte Arbeitende vor Computerterminals sitzen und die Fotos, die von Usern weltweit auf Facebook hochgeladen werden, auf verbotene Inhalte hin klassifizieren müssen, sind ein besonders brutales Beispiel für die Einbindung einer menschlichen kognitiven Kapazität in ein Computernetzwerk.²⁸ Es markiert ein extremes Ende des Spektrums dessen, was als *Human-Aided AI* bezeichnet werden kann, weil in diesem Fall eine globale wirtschaftliche Ungleichheitssituation die Rahmenbedingungen dafür bietet, Arbeitskräfte explizit und unter schweren gesundheitlichen Folgen für eine kognitive Fähigkeit auszuheben, die die automatische Bilderkennung noch nicht ganz ersetzen kann. Automatische Bilderkennung sortiert das Bildmaterial nämlich lediglich vor, so dass nur die wirklich harten und Zweifelsfälle den Klickarbeiter_innen auf den Philippinen vorgelegt werden, was deren psychische Belastung bis hin zur Posttraumatischen Belastungsstörung noch erhöht, weil damit die Dichte brutaler und schwer verarbeitbarer Bildinhalte im Stream der anzuschauenden und zu klassifizierenden Items steigt.

Zur Thematisierung von Entmündigung eignet sich ein schlichteres und weniger schmerzhaftes, dafür in der täglichen Interaktion zwischen Benutzer und Maschine lokalisiertes Beispiel für *Human-Aided AI*. Im Design der Facebook-Benutzeroberfläche hat es irgendwann zwischen dem Jahr 2009 und 2013 eine Umstellung gegeben. Vor dieser Umstellung enthielt die »Info«-Sektion des eigenen Profils verschiedene Felder wie »Interests«, »Favorite Music«, »Business Skills« etc., die jeweils eine freie Eingabe von Text erlaubten. So speicherte eine

28 | Siehe Reuter 27.04.2016 sowie Der Standard 29.04.2016.

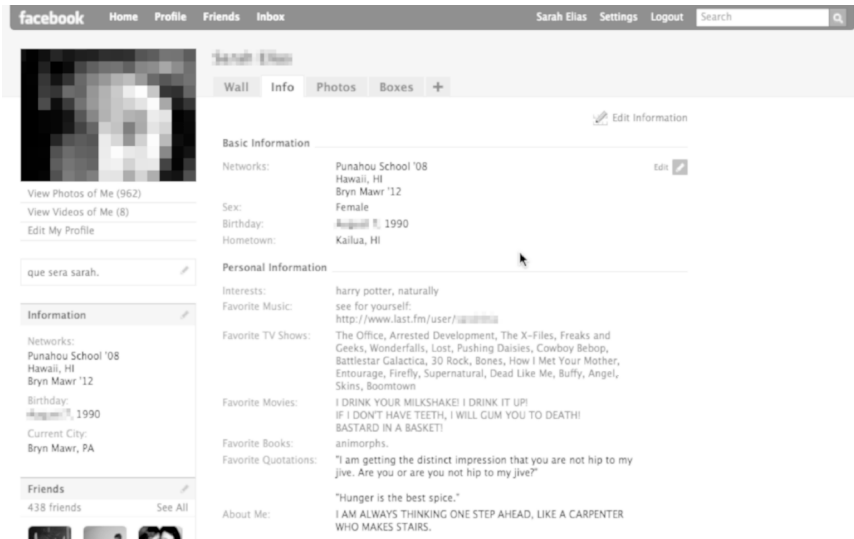


Abbildung 2: Facebook-Profil im Jahr 2009 (Screenshot vom 25.04.2009). Ungepars- te Informationen im Stil eines Steckbriefs. Quelle: Sarah Elias, <http://gandt.blogs.bryn mawr.edu/web-papers/web-papers-4-multimedia-projects/facebook-mosaic/>. Verpixelun- gen hinzugefügt.

Nutzerin im Feld »Interests« zum Beispiel wörtlich: »harry potter, naturally« (siehe Abbildung 2). Nach der Umstellung wurde die Möglichkeit der freie Eingabe von Informationen in den meisten Feldern eines Facebook-Profiles abgeschafft und durch einen partizipativen Echtzeit-Parsingmechanismus ersetzt. »Parsen«, das ist in der Informatik das Problem der strukturellen und semantischen Zergliederung und maschinellen Erfassung einer Dateneingabe. Gibt man heute in einem der Felder eines Facebook-Profiles – im Folgenden am Feld »Professional Skills« demonstriert, siehe 3 – etwas ein, dann wird noch während des Tippens direkt unterhalb des Eingabefeldes eine ständig aktualisierte (inkrementelle) Suchre- sultatliste angezeigt, die zu dem eingegebenen Wortfragment hinterlegte Items anzeigt. Es ist nicht möglich, ein Wort als »Professional Skill« einzutragen, das nicht aus dieser Liste gewählt wird; es handelt sich um eine Hybridfunktion aus freier Eingabemöglichkeit und Auswahl aus einem festgelegten Verzeichnis von Optionen. Aus Sicht des Betreibers ist der zentrale Gewinn dieser Technik, dass damit das Problem der strukturellen Aufarbeitung (Parsing) der Eingabe vermieden wird. Eine trickreiche Gestaltung der Mensch-Maschine-Interaktion greift in- teraktiv in den Prozess der Dateneingabe ein, so dass die Nutzer_in dabei selbst die Aufgabe übernimmt, ihre Angaben einer *serverseitig registrierten semantischen Kategorie zuzuordnen*. Es wird damit ein strukturelles Item in der Datenbank ge- speichert, die erfasste Information ist sofort in wirtschaftlich verwertbarer Weise

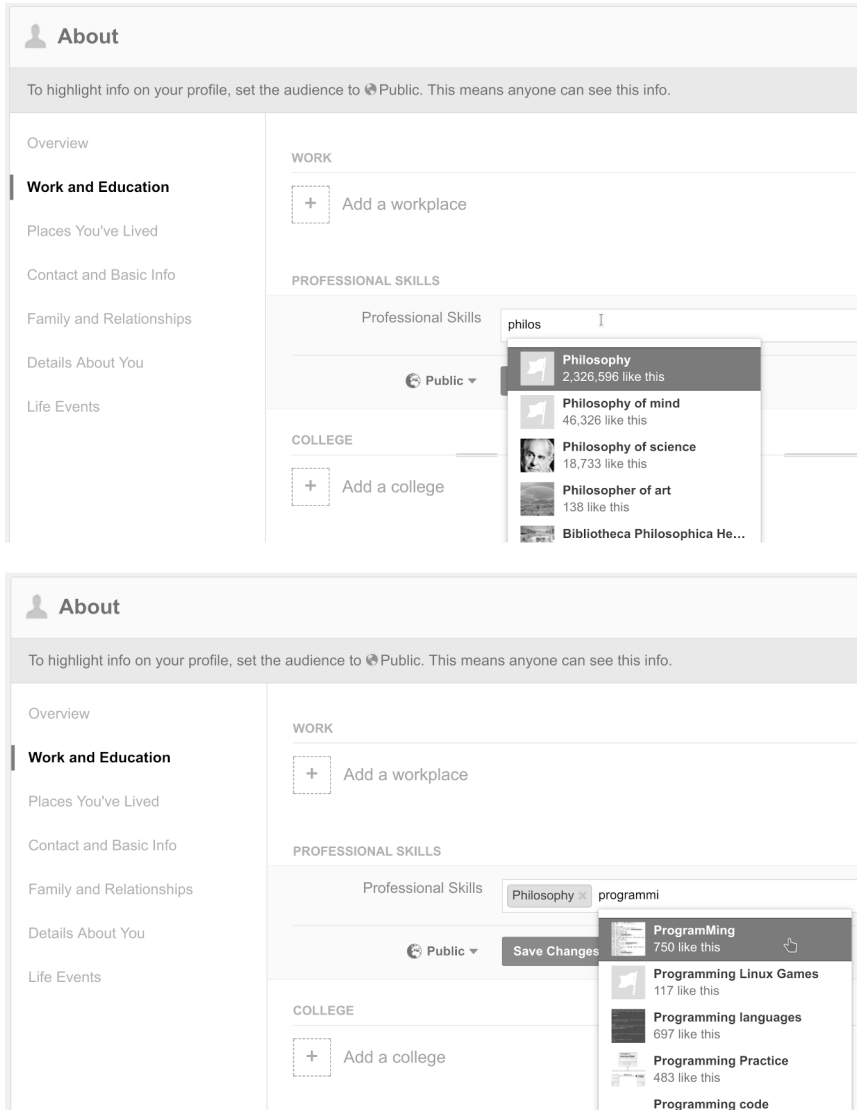


Abbildung 3: Facebook-Profil im Jahr 2018. Das Parsing-Problem wird umgangen, indem die Nutzer_in interaktiv dazu genudged wird, semantisch aufgearbeitete Daten einzugeben.
Quelle: facebook.com, Screenshots des Verfassers vom 30.01.2018.

aufbereitet. Eingaben wie »harry potter, naturally« oder »see for yourself: <http://www.last.fm/user/...>« (Fig. 2) sind in dieser technischen Rahmung der Nutzerinteraktion nicht möglich – sie sind aber auch nicht erwünscht, weil sie sich nicht automatisiert verwerten lassen.

Es handelt sich hierbei um eine gezielte und auf dem Weg des Interface-Designs implizit erwirkte Einbindung der menschlichen kognitiven Fähigkeit in ein dezentrales Computersystem, mit dem Ziel, strukturell und semantisch verwertbare Informationen anstatt bloß einer schwer interpretierbaren, eventuell frei gestalteten Zeichenkette als Eingabe zu erhalten. Nach dem kybernetischen Prinzip einer Feedbackschleife wird es durch dieses Interaktionsdesign der Nutzer_in überlassen, zu ermitteln, welche der vorgeschlagenen Kategorien am ehesten der semantischen Interpretation der von ihr eingegebenen Zeichenkette entspricht – anstatt dass eine Maschine die Nutzereingabe im Hintergrund interpretieren müsste und dabei Fehler machen würde, die nie korrigiert werden können, wenn das Resultat der Nutzer_in nicht sofort angezeigt wird. Der technisch aufwändige und fehleranfällige Schritt der semantischen Interpretation oder Aufbereitung (Parsing) von Nutzerdaten wird hiermit also direkt an die Nutzer_innen outgesourced, als wären sie kleine informationsverarbeitende Module (»wet-ware«) in einem großen Computernetzwerk.

Diese strukturelle und damit maschinenverarbeitbare Erfassung von Profildaten macht es der Plattform erst möglich, zum Beispiel die persönlichen Interessen und beruflichen Fähigkeiten der Nutzer_innen als Datenpunkte in den »sozialen Graphen« einzufügen. Nutzer_innen können dadurch nach ihren Gemeinsamkeiten in diesen Dimensionen verknüpft werden und durch den Echtzeit-Parsingmechanismus werden sie überdies dazu aufgefordert, diese Verknüpfungen *selbst* auszuweisen. Bei den auswählbaren Items der Liste handelt es sich nicht um einen von Facebook-Mitarbeiter_innen ausgearbeiteten Katalog, sondern um all das, was andere Benutzer auf Facebook »geliked« haben – die Liste ist also nicht starr codiert, sondern wird nach einem dynamischen und *immanenten* Mechanismus der Plattform generiert. Und umgekehrt: Eine Information über den User, die diesen nicht mit anderen Usern verknüpft, ist für Facebook wertlos und deshalb ihre Erfassung uninteressant – daher kommt es, ökonomisch betrachtet, nicht auf freie Eingabemöglichkeiten an. Nur durch strukturierte Informationen erhöht sich die verfügbare Datenauflösung für die Zwecke automatisierter Weiterverarbeitung, etwa zur Gewinnung eines psychometrischen Persönlichkeitsbildes der Nutzer_in, das im Rahmen von Risikomanagement-Services, *targeted ads* und anderen Big-Data-getriebenen Analysen vermarktet werden kann und wird.

3 Digitale Entmündigung

In einem reichhaltigen Sinn ist eine Entscheidung freiwillig, wenn sich ein Subjekt bewusst ist, worüber es entscheidet, die Sache ihm also nicht untergeschoben wird; wenn das Subjekt hinreichend informiert und gebildet ist, um über die Sache Bescheid wissen und urteilen zu können; wenn das Subjekt auch über veritable Alternativen zu der Entscheidung verfügt, wenn die Situation also eine *echte* Wahlmöglichkeit bietet. Nach diesem Verständnis ist die *Consent-or-leave-Policy* bei Nutzungsbedingungen – zustimmen oder nicht nutzen – keine freie Entscheidungssituation. In dieser Differenzierung deutet sich ein starker und normativer Begriff der Freiwilligkeit an, der nicht nur das entscheidende Subjekt, sondern die Rahmung der Entscheidung betrifft. Denn diese Rahmung kann Freiwilligkeit ermöglichen oder strategisch sabotieren. Wenn Wahlarchitekturen digitaler Interfaces das Einsatzfeld einer Macht- und Marktstrategie bilden, betrifft diese Norm besonders auch die Form der *Ansprache* (durch das technische Gerät), die ein Subjekt zur Freiwilligkeit ermächtigen kann oder nicht.

Dieser Aspekt bleibt radikal unbeleuchtet in einer Diskussion um Datenschutz, die sich zu sehr auf illegal erlangte Daten fokussiert. Datenschutz wird dadurch verengt auf die Gefahr des personalisierten Angriffs – jemand interessiert sich für *meine* Daten. Was jedoch gesellschaftlich und sozial viel schwerer wiegt ist nicht der Einbruch in die Geheimnisse des Einzelnen *als Einzelnen*, also die Erlangung dessen, was jemand Konkretes nicht preisgeben wollte. Sondern es sind die Wahrscheinlichkeitsaussagen, die man regulär über jedes *beliebige* Individuum anhand eines Massendatensatzes treffen kann. Die Gefahr bilden die Daten, die Nutzer_innen täglich freiwillig zur Verfügung stellen, und die *abgeleiteten Daten* (Korrelationen mit anderen Usern), die daraus generiert werden.²⁹ Die heute öffentlich geführte Debatte um Datenschutz hingegen fokussiert auf einen liberalistischen Individualismus und verliert damit die fundamentalen Transformationen des Sozialen und Politischen aus den Augen, die die ökonomische, polizeiliche und politische Verwendung von Daten als *Massendaten* möglich macht.

Der Massendatensatz fällt durch die *kollektive* Nutzung von Plattformservices an und erlaubt es, *beliebige* Individuen umfangreich und mit hoher Auflösung in Relation zu anderen Individuen einzuordnen, zu beurteilen und zu diskriminieren. Diese Effekte werden nicht durch die möglichst hohe Detailtiefe der über ein *bestimmtes* Individuum erhobenen Daten möglich, sondern durch die Menge der Vergleichsobjekte und Vergleichsdaten. Man trägt zu den sozialen Selektionseffekten, zu ökonomischen, politischen und sozialen Ausgrenzungen und Hierarchisierungen durch Big Data auch dann bei, wenn man selbst durch die Erhebung seiner Daten keine negativen Effekte zu befürchten hat und *als Einzelner*

29 | Siehe zur Verwendung abgeleiteter Daten in einer Vielzahl gesellschaftlicher und ökonomischer Bereiche ausführlich O’Neil 2016 und speziell zu den sozialen Implikationen auch die Kampagne <http://www.socialcooling.com/> von Tijmen Schep.

kein Problem darin sieht (vgl. O’Neil 2016). Die Subjektivität derer, die alltäglich ihre Daten zur Verfügung stellen, »weil es so bequem ist«, weil sie »ja nichts zu verbergen haben« oder »es sowieso schon zu spät« sei, ist aus diesem Grund ein politisches und gesellschaftliches Problem. Das Phänomen der »freiwilligen« Datenpreisgabe muss viel genauer untersucht und auch auf Ebene der zugrundeliegenden subjektiven Mechanismen kritisiert werden.

Mit dem Begriff »Subjektivität« in diesem Zusammenhang ist gemeint, dass das hier besprochene Nutzerverhalten selbst in einer Relation wechselseitiger Hervorbringung mit den technischen Dispositiven steht. Nutzer_innen sind den medialen Oberflächen, dem Interfacedesign, in denen die oben besprochenen Nudging-Techniken am Werk sind, täglich ausgesetzt. Es gibt so etwas wie eine Techniksубјektivierung in digitalen Räumen, das ist der über längere Zeiträume sich einstellende Effekt der Hervorbringung einer konkreten Art und Weise, technische Services und Interfaces wahrzunehmen und sich in ihnen und zu ihnen – und vermittelt ihrer auch zu anderen und zur Gesellschaft – zu verhalten. Teil dieser Techniksубјektivierung ist eine geschärfte Wahrnehmung dafür, was eine »gute« und was eine »schlechte« Benutzerführung ist. Darunter fällt auch eine Aversion gegen Konfrontationen mit technischen Details; das Paradigma des »user-centered designs«³⁰ formuliert genau die Anspruchshaltung, dass die Technik benutzbar sein muss, ohne von der Nutzer_in zu verlangen, einen Informatikabschluss zu haben, ein Manual zu lesen, überhaupt nachdenken zu müssen.³¹

In einem Zeitalter jedoch, wo der Diskurs des Designs technischer Interfaces es methodologisch genau darauf abstellt, menschliche Regungen – affektiv, psychologisch, körperlich, kognitiv, sozial, politisch – detailliert quantitativ erfassen und antizipieren zu können, ist gerade die Kategorie des »Willens«, die ja im Begriff der »Freiwilligkeit« vorkommt, eine hart umfochtene und äußerst unzuverlässige Kategorie. Das Wissen dieses Diskurses möchte nicht nur Interfaces gestalten, in denen *im subjektiven Gefühl der Freiwilligkeit* mit maximaler Wahrscheinlichkeit eine bestimmte, von anderen vorgefasste Entscheidung getroffen wird. Das Momentum dieser Entwicklung reicht viel weiter, bis zur Vision etwa, dass die Suchmaschine Google auf Grundlage prädiktiver Analysen Antworten liefert, noch bevor überhaupt eine Frage gestellt werden muss – genauso wie Amazon sich im Jahr 2013 ein Verfahren des »pre-shipping« patentieren ließ, nach dem auf Grundlage voraus kalkulierten Kundenverhaltens Produkte zum Kunden nach Hause oder in seine Nähe geliefert werden, noch bevor eine Bestellung eingegangen ist.³²

30 | Siehe Donald Norman (1988): *The Design of Everyday Things*, eines der Manifeste zum *user-centered design*, das zuerst unter dem Titel *The Psychology of Everyday Things* erschien.

31 | *Don’t Make Me Think!* ist der Titel eines in den 2000er Jahren beliebten Web-Usability-Standardwerks von Steve Krug (2005), worin er das gleichlautende Prinzip auch zur Regel Nr. 1 für gutes Webdesign erklärt.

32 | US patent #US008615473, vgl. <http://techcrunch.com/2014/01/18/amazon-pre-ships/>.

Das Paradigma des *user-centered designs*, die Erwartung von Einfachheit und Intuitivität der Bedienung, ist in Bezug auf die Frage der Techniksubjektivierung einer der prägenden Trends in der aktuellen Dekade. Es handelt sich bei der Bemühung um Nutzerfreundlichkeit nicht um ein spätes sozialkompetentes Erwachen von Computer-Nerds, sondern um eine Machtstrategie, um einen Willen zur Macht, der sich im Design materialisiert, und um ein Bestreben, Design-Hegemonien in Bezug auf die Interaktionsschemata mit Technologien zu etablieren, wie sich bereits im Studium der *Developer Guidelines* des Facebook Logins andeutete. Zentral für diese Machtstrategie, wenn man sie als Strategie der Subjektivierung von Nutzer_innen untersuchen möchte, ist der Stil der *Ansprache* der Nutzer_innen zum Beispiel durch *landing pages*, Dialogboxen, Allgemeine Geschäftsbedingungen und Interfaces. Diese Ansprache wird mit viel Aufwand so eingerichtet, dass sie ihre Nutzer_innen entmündigt, wenn Entmündigung bedeutet, (1) den Nutzer_innen nichts zuzutrauen, also rein intuitives und ›bequemes‹ Entscheidungsverhalten von ihnen zu erwarten, und (2) ihnen keine uneindeutigen Optionenräume zu überlassen, die sie vor Gabelungspunkte stellen, an denen sie eventuell vor lauter Wahlmöglichkeiten nicht weiterkommen.

Das »UX Design« macht dem User die Sachen also nicht deshalb einfach, weil es ihm entgegenkommen möchte, sondern um ihn einzuhegen. Subjektiv korreliert das nicht nur mit der steigenden Bereitschaft, sondern mit der zur sozialen *Norm* gewordenen Resignations- und Ohnmachtshaltung gegenüber Technik, in der gefordert wird, das Technische am Technischen möglichst nicht sehen zu wollen, es hinter einer »streamlined«, fließenden, erlebnisreichen Bedienoberfläche verkapselt zu wissen – Grundkonzept der Marke Apple seit den 1980ern. Die Sachen zu *verkomplizieren*, ihre Details und Ambivalenzen sichtbar zu machen, ist im kulturellen und subjektiven Verhältnis zu Technik verpönter denn je – und in diesem Punkt liegt eine unbemerkte Komplizenschaft einer über politische, soziale und Klassengrenzen hinweg weit verbreiteten subjektiven Einstellung mit den ökonomischen Interessen von GAFA.

Literatur

- Foucault, Michel (1978 [1977]). »Ein Spiel um die Psychoanalyse. [Interview]«. In: *Dispositive der Macht. Michel Foucault über Sexualität, Wissen und Wahrheit*. Berlin: Merve, S. 118–175.
- Krug, Steve (2005). *Don't Make Me Think! A Common Sense Approach to Web Usability*. 2. Aufl. Berkeley: New Riders.
- Mühlhoff, Rainer (2018). »Digitale Entmündigung und ›User Experience Design‹. Wie digitale Geräte uns nudgen, tracken und zur Unwissenheit erziehen«. In: *Leviathan – Journal of Social Sciences* 46 (4).

- Mühlhoff, Rainer (2019). »Menschengestützte Künstliche Intelligenz: Über die sozial-medialen Voraussetzungen von Deep Learning«. In: *Zeitschrift für Medienwissenschaft*. Jg. 11, Heft 21. Im Druck.
- O'Neil, Cathy (2016). *Weapons of Math Destruction. How Big Data Increases Inequality and Threatens Democracy*. Penguin.
- Sajja, Priti Srinivas und Akerkar, Rajendra (2012). *Intelligent Technologies for Web Applications*. Boca Raton, London und New York: CRC Press.
- Thaler, Richard H. und Sunstein, Cass R. (2008). *Nudge: Improving decisions about health, wealth, and happiness*. New Haven: Yale University Press.

Online-Quellen

- Charkham, Avi (25.08.2012). *5 Design Tricks Facebook Uses To Affect Your Privacy Decisions*. URL: <https://techcrunch.com/2012/08/25/5-design-tricks-facebook-uses-to-affect-your-privacy-decisions/>.
- Der Standard (29.04.2016). »Verspeiste Fäkalien, Sex mit Tieren«: Facebook-Zensoren leiden unter Horrorbildern. URL: <http://derstandard.at/2000035900517/Verspeiste-Faekalien-Sex-mit-Tieren-Facebook-Zensoren-leiden-unter-Horrorbildern>.
- Kelly, David (02.01.2014). *How to Use the Information Inside Google's Ved Parameter*. URL: <https://moz.com/blog/inside-googles-ved-parameter>.
- Mediative (2014). *The Evolution of Google's Search Results Pages & Effects on User Behaviour*. URL: <http://www.mediative.com/whitepaper-the-evolution-of-googles-search-results-pages-effects-on-user-behaviour/>.
- Ny, Ken (08.06.2016). *Wie funktioniert das click tracking bei Google Search?* URL: <https://www.blocko11.de/2015/06/08/wie-funktioniert-das-click-tracking-bei-google-search/>.
- Oxley, Ian (20.03.2014). *New HTML5 Attributes for Hyperlinks: download, media, and ping*. URL: <https://www.sitepoint.com/new-html5-attributes-hyperlinks-download-media-ping/>.
- Pichsenmeister, David (02.12.2016). URL: <https://venturebeat.com/2016/12/02/moving-from-ai-assisted-humans-to-human-assisted-ai/>.
- Resnik, Tim (22.05.2013). *Decoding Google's Referral String (or, how I survived Secure Search)*. URL: <https://moz.com/blog/decoding-googles-referral-string-or-how-i-survived-secure-search>.
- Reuter, Markus (27.04.2016). *Die digitale Müllabfuhr: Kommerzielle Inhaltsmoderation auf den Philippinen*. URL: <https://netzpolitik.org/2016/die-digitale-muellabfuhr-kommerzielle-inhaltsmoderation-auf-den-philippinen/>.
- sshay77 (18.07.2015). *google-search-url-parameters-query-string*. GitHub Gist Repository. URL: <https://gist.github.com/sshay77/4b1f6616a7afabc2a>.
- Wittersheim, Aaron (31.03.2016). *The Approaching Darkness: The Google Referral URL In 2016*. URL: <https://www.straightnorth.com/insights/approaching-darkness-google-referral-url-2016/>.