

II.2.5 Datafizierung, Algorithmizität und Künstliche Intelligenz

Von Menschen und Maschinen werden aufgrund immer kleinerer Sensoren immer ungeheurere Mengen an Daten produziert. Jede Bewegung und Tätigkeit kann in Daten überführt werden und wird damit für Computer lesbar (maschinenlesbar), bearbeitbar und auswertbar. Dabei nimmt nicht nur die Menge an erhobenen Daten zu, sondern auch die Bandbreite an unterschiedlichen Daten. Zugleich führt die technische Entwicklung zum einen zu immer größeren und günstigeren Speichermedien, die ein prinzipiell unbegrenztes Vorhalten aller Daten ermöglichen (wenn dem nicht Datenschutzgesetze entgegenstehen). Zum anderen ermöglicht die steigende Prozessorleistung eine massenhafte oder immer ausgeklügeltere Verarbeitung und Auswertung der vorhandenen Daten. Zusammengefasst findet damit eine »Datafizierung der Welt« statt (Ramge 2019a: 49).⁴⁴

Bezogen auf soziales Handeln und gesellschaftliche Praktiken, sprechen unter anderem Prietl und Houben (2018: 7) davon, dass »soziotechnische Prozesse der *Datafizierung* [Herv. i. O.] des Sozialen im Zentrum gegenwärtiger gesellschaftlicher Transformationen stehen.« So »vollzieht sich das Soziale zunehmend in numerisierten Umwelten« (ebd.: 10). Im Kontext des sozialen und gesellschaftlichen Lebens, ebenso wie des wirtschaftlichen und politischen Handelns fallen nicht nur immer mehr Daten an, sondern werden auch gesammelt und ausgewertet (vgl. Dencik et al. 2019: 3). Harari (2017: 497) spricht von dem »Dataismus« – einer »Datenreligion«, der zufolge das »Universum aus Datenströmen« bestehe und der Wert von allem und jedem an seinem »Beitrag zur Datenverarbeitung« gemessen werde.

Datenerfassung erfolgt nicht mehr auf bestimmte Zeitpunkte und Ziele beschränkt, sondern Daten werden automatisch sowohl als Haupt- als auch Nebenprodukt bei der Nutzung (informations-)technischer Systeme generiert. Daten werden so nicht mehr nur aktiv von Dritten erfragt, sondern auch selbst erhoben (wofür insbesondere die *Quantified Self-Bewegung* [Selbstvermessung] steht) und bereitgestellt (soziale Netzwerke).⁴⁵

Wenn Daten im digitalen Zeitalter eine entscheidende Rolle zukommt, dann bringen die Verfügbarkeit und der Besitz von Daten Machtpotenziale mit sich. Damit geht

44 Wobei Gugerli (2018: 193) darauf hinweist, dass es bereits am Beginn der Computerära mehr Zeit brauchte, um die »Welt maschinenlesbar zu machen [...], als Daten zu verarbeiten.«

45 Die verschiedenen Quellen aus denen Daten stammen, beziehungsweise die unterschiedlichen Stellen an denen Daten erzeugt oder erhoben werden, führen unter anderem zu der Debatte um Dateneigentum. Verdeutlicht wird die dahinterstehende Problematik etwa an einer Autofahrt, bei der Sensoren im Auto beispielsweise Daten über die gefahrene Strecke, das Bremsverhalten oder die Geschwindigkeit erfassen. Wem gehören diese Daten? Der Eigentümerin des Autos? Dem Fahrer des Autos? Dem Hersteller des Autos? Oder der Versicherung, die von der Autobesitzerin gegen Vergünstigung der Versicherungspolice eine zusätzliche Schnittstelle zur Datenweitergabe oder einen Sensor zur Datenerfassung hat installieren lassen?

Die hierüber geführten Debatten werden auch nicht dadurch einfacher, dass Daten verlustfrei dupliziert werden können und damit ihr Besitz nicht wie bei einem physischen Gut exklusiv sein muss. Erschwerend kommen Datenschutzfragen hinzu (etwa inwieweit ein Autobesitzer die Autofahrerin über die Datenerhebung informieren muss) – insbesondere, wenn es um noch sensiblere Daten wie etwa Gesundheitsdaten geht.

eine zentrale Frage einher: Wer soll Zugriff auf diese Daten haben, und welcher Umgang mit ihnen ist erwünscht? Die Antwort auf diese Frage hängt unter anderem stark damit zusammen, von wem die Daten erhoben, gespeichert und ausgewertet werden – und betrifft damit die Themenfelder Datenschutz und Privatsphäre.⁴⁶ Aus der ungleichen Verteilung von Daten gehen daher asymmetrische Machtrelationen und Machtbeziehungen hervor (vgl. Houben/Prietzl 2018b: 348ff.). Besonders deutlich werden diese in der (westlichen) Datenkonzentration bei wenigen großen (US-amerikanischen) Internetkonzernen. Die bisherige Entwicklung digitaler Plattformen und Internetkonzernen zeigt bislang in eine eindeutige Richtung: »[O]hne neue kartellrechtliche Gegenmaßnahmen führen sie langfristig fast unweigerlich zu Datenmonopolen« (Ramge 2019b: 52). Bei deren aufgrund des Skalen- und Netzwerkeffekts zur Monopolbildung⁴⁷ tendierenden digitalen Plattformen fallen erhebliche Datenmengen an. Dienstleistungen werden formal gratis angeboten, in Wirklichkeit aber über expansive Datenausnutzung bezahlt. Darüber hinaus macht »Partizipation« in der Plattformökonomie Nutzer:innen zu Prosumer:innen (Producer:innen und Consumer:innen zugleich), wodurch diese ohne Entlohnung zur Weiterentwicklung von Produkten (kreieren eigener Designs) oder zur Bewertung und Empfehlung von Waren und Dienstleistungen (Reviews, Ratings, Tests) beitragen (vgl. Stalder 2017: 66f.). Diese Monopolisierung von Daten stellt eine zentrale Gefahr dar, für die Zuboff (2018) den Begriff des Überwachungskapitalismus [*surveillance capitalism*] geprägt hat (siehe hierzu auch Kapitel IV.4.1). Die Machtasymmetrie zwischen den Nutzer:innen und den Anbietern von Internetdiensten ist eine der wahrgenommenen Problemdimensionen, die einen Ausgangspunkt für die Suche nach einer angemessenen staatlichen Regulierung darstellt, was an späterer Stelle noch ausführlich behandelt wird (siehe Kapitel V.3).

Aber auch für staatliche Steuerung direkt spielen die (Nicht-)Verfügbarkeit von und der Zugang zu Daten (als zentrale Ressource für Steuerungswissen) eine zentrale Rolle. Wo etwa früher staatliche Stellen aufgrund ihrer (bevölkerungs-)statistischen Erhebungen, etwa durch die Statistikämter auf Landes- und Bundesebene, die »Souveränität über Datenregime« innehatten, findet heute die Datensammlung in erheblichem Umfang im privatwirtschaftlichen Sektor statt (Prietzl/Houben 2018: 8, siehe auch 10).

»Woher der Staat sein Wissen bezieht – nenne man dies Herrschaftswissen, Regierungswissen oder Regelungswissen – ist im Staat der Wissensgesellschaft nicht beliebig, sondern ein zentrales Governanceproblem« (Schuppert 2013: 45).

Noch weiter geht Pistor (2021: 4), wenn sie die Diskussion um die Staatlichkeit im digitalen Zeitalter und den digitalen Staat gleich ganz auf die Datenfrage verkürzt:

46 Von vielen Unternehmen, die Smart Devices zur Datenerhebung vertreiben, wird diese Frage auf der technischen Ebene relativ eindeutig beantwortet: Unabhängig davon, ob das Gerät selber zur Speicherung oder sogar Auswertung der Daten in der Lage wäre, findet häufig beides in der Cloud – und damit unter direkter Zugriffsmöglichkeit – des Herstellers statt.

47 Aus einer nicht auf den Westen zentrierten Perspektive bilden sich im Internet separierte Oligopole mit insbesondere jeweils englisch-, russisch- und chinesischsprachigem Einzugsbereich.

»Die neuen digitalen ›Staaten‹ könnten in erster Linie auf Daten und nicht auf einem Staatsgebiet basieren [...]. Wenn diese Hypothese Staatlichkeit im digitalen Zeitalter zumindest annähernd richtig erfasst, dann wird darunter eine andere Art von Staat und ein neuer Typus des Souveräns zu verstehen sein: der Datencontroller als neuer Datensouverän.«

Auch wenn man dieser eingeschränkten Perspektive nicht folgt, ist es nicht verwunderlich, dass im digitalen Zeitalter die Gestaltung des Zugangs zu Steuerungswissen für den Staat zentral – gleichzeitig aber häufig nicht unumstritten – ist. Erkennen lässt sich dies zum einen an Debatten um die Gefahr einer Expertokratie aufgrund der Bedeutung von externer Expertise, Sachverständigen, Gutachten, Anhörungen und Kommissionen (siehe Kapitel VI.2.2.1). Zum anderen werden die Versuche, vorhandene Daten zusammenzuführen und für Staat, Wissenschaft und Wirtschaft nutzbar zu machen (siehe etwa Kapitel V.1.3.4 zum Umgang mit Gesundheitsdaten), kritisch begleitet, unter anderem, weil Datenmonopole und damit Überwachungskapazitäten durchaus auch bei staatlichen Stellen entstehen könnten (siehe Kapitel IV.4.1).

Aus der Datafizierung ergibt sich die Algorithmizität der modernen, digitalen Welt. Das Vorhandensein mannigfaltiger Daten (Big Data) ermöglicht erst deren automatische Auswertung und benötigt gleichzeitig die algorithmische Analyse.

»Auf dem Weg in das ›Datenzeitalter‹ erkennen wir digitale Algorithmen als notwendiges Instrumentarium, um das exponentielle Wachstum an Daten bewältigen zu können, also überhaupt handhabbar, navigierbar, verwaltbar zu machen« (Mohabbat-Kar/Parycek 2018: 9).

Das heißt, die Welt des digitalen Zeitalters »ist geprägt durch automatisierte Entscheidungsverfahren, die den Informationsüberfluss reduzieren und formen, sodass sich aus den von Maschinen produzierten Datenmengen Informationen gewinnen lassen, die der menschlichen Wahrnehmung zugänglich sind und zur Grundlage des singulären und gemeinschaftlichen Handelns werden können« (Stalder 2017: 13). Dahinter verbirgt sich zum einen die Sehnsucht nach der Berechenbarkeit der Welt – und damit auch ihrer weiteren Entwicklung. Die Zukunft ist dann »nicht mehr der offene Horizont, sondern die nichtkontingente, verrechnete Wahrscheinlichkeit« (Feustel 2018: 152).⁴⁸ Die Datafizierung soll Sicherheit schaffen, indem der Raum des Ungewissen verkleinert wird. Exakte Daten, umfassendes Wissen und die Quantifizierung kleinster Details suggerieren Eindeutigkeit und Wahrheit (vgl. Lotter 2019: 44; Bauer 2018: 38). Dass mit zunehmendem Informationsumfang vielfach eher das Gegenteil eintritt, wird in Kapitel II.2.8 vertieft. An dieser Stelle soll grundlegend auf die Wirkungsweise von Algorithmen eingegangen werden, wozu es zunächst genügt festzuhalten, dass es auch im Angesicht von Datafizierung und Quantifizierung nicht um Sicherheiten (im Sinne einer eindeutigen Zukunftsvorhersage), sondern um Möglichkeiten (und damit Wahrscheinlichkeiten) geht – den

48 Weitergedacht verbirgt sich dahinter »das auf die Spitze getriebene Newton'sche Weltbild: Wer alle natürlichen Gesetze und Kräfte kennt und über vollständiges Wissen sämtlicher Zustände aller Materieteilchen zu einem bestimmten Zeitpunkt verfügt, kann die Zukunft – den Zustand der Materieteilchen zu einem späteren Zeitpunkt – akkurat vorausberechnen« (Mamczak 2014: 13).

»Zukunftswissen ist [...] immer probabilistisches Wissen« (Mamczak 2014: 14). Genau um so eine Berechnung der Zukunft geht es auch häufig beim Einsatz von Algorithmen.

In einer Vorstufe dienen Algorithmen aber auch dazu, neue Daten zu generieren. Durch die algorithmische Verknüpfung von Daten entstehen neue Daten. Unter dem Begriff des *Data Mining* [Datengewinnung] wird sogar explizit das automatische (durch Algorithmen) Finden von Mustern und Zusammenhängen in Big Data verstanden.⁴⁹ Diese sollten nicht mit neuem Wissen verwechselt oder gleichgesetzt werden. »Wissen setzt Bewusstsein voraus; Bewusstsein haben nur Lebewesen. Computer können nicht denken und nicht fühlen; allenfalls können sie Wissen, Denkprozesse und Gefühle simulieren« (Bull 2019: 68). Neben der Datengenerierung dienen Algorithmen der Entscheidungsfindung. Aus der Perspektive der Berechenbarkeit unterstützen Algorithmen mit ihrem Ergebnis entweder die Entscheidungsfindung von Menschen⁵⁰ oder entscheiden selbst – in dem Sinn, dass sie auf Basis ihres Ergebnisses eine Folgeaktion auslösen.⁵¹ Auch die zugrunde liegenden Berechnungen erzeugen dabei kein Wissen. Wenn ein Algorithmus bei der Kreditanfrage eines Kunden diesen ablehnt, dann hat das nichts damit zu tun, dass der Algorithmus aus den vorliegenden Daten neues Wissen über diesen Kunden generiert hätte.⁵² Seine Entscheidung basiert nicht auf Wissen, sondern auf Wahrscheinlichkeiten. Und die Wahrscheinlichkeit, dass der Kunde seinen Kredit (nicht) zurückzahlen können wird, basiert auf der Ähnlichkeit der vorliegenden Daten des Kunden zu den Ausprägungen bekannter Fälle, in denen ein Kredit (nicht) zurückgezahlt werden konnte. Ein zentrales Problem bei einer solchen algorithmischen Berechnung von Entscheidungen: sie stellen nur vermeintliche Eindeutigkeit angesichts eines mehrdeutigen Phänomens her, weil sie, »egal wie wahrscheinlich, immer nur eine unter verschiedenen Möglichkeiten in den Rang kommender Wirklichkeit erhebt« (Feustel 2018: 153). Man könnte auch sagen: Es handelt sich um eine Wette auf die Zukunft.

Hinter dieser Wette stehen Algorithmen, bei denen es sich um wenig mehr als mathematisch-statistische Verfahren handelt. Daran ändert auch der aktuelle Peak in der Forschung zu Künstlicher Intelligenz (KI) [*Artificial Intelligence*; *AI*] nicht grundlegend etwas. Der Begriff der Intelligenz führt hier aufgrund seiner umgangssprachlichen Bedeutung in die Irre, weshalb der Begriff KI auch bei KI-Forschern nicht unumstritten ist. Zumindest sollte zwischen starker KI [*strong AI* oder *general AI*] und schwacher KI [*weak AI* oder *narrow AI*] unterschieden werden. Eine starke KI wäre durch intelligentes Verhalten sowie die Fähigkeit zu logischem Denken und zu Kommunikation in natürlicher Sprache gekennzeichnet, sodass sie unterschiedlichste kognitive Anforderungen

49 Hierbei geht es also gerade nicht darum, Antworten auf gestellte Fragen zu bekommen, also Wissen zu generieren. Der Algorithmus findet hier vielmehr explorativ Antworten, ohne dass man die Frage kennen würde – insofern handelt es sich um Daten, die erst analysiert, interpretiert und mit Kontext versehen werden müssen, bevor sie neues Wissen darstellen. Der um die letzteren Bearbeitungsschritte erweiterte Prozess wird daher auch als *Knowledge Discovery* bezeichnet.

50 Die Algorithmen der Navigationssoftware unterstützen so etwa bei der Auswahl der kürzesten, schnellsten oder landschaftlich schönsten Wegstrecke.

51 Dies ist etwa der Fall, wenn in der Finanzindustrie ein Algorithmus automatisch je nach Kursentwicklung und berechneter Prognose eigenständig Aktien kauft oder verkauft.

52 Algorithmen liefern kein Abbild der Realität, sondern konstruieren diese mit (Stalder 2017: 194).

erfüllt (vgl. Spath 2018: 506).⁵³ Eine solche »Starke KI ist bis auf Weiteres Science-Fiction« (Ramge 2019b: 19). Stand der Technik ist also (auch auf absehbare Zeit) die schwache KI mit Anwendungen etwa in der Text-, Bild- und Spracherkennung oder bei Experten-, Assistenz- und Navigationssystemen. Es handelt sich um jeweils hochspezielle Systeme, in denen die KI eine ganz spezifische Funktion oder Aufgabe übernimmt. Die KI-Forschung setzt auf maschinelles Lernen [*machine learning*], selbstlernende Algorithmen und *deep learning* [mehrschichtiges Lernen] in neuronalen Netzen.⁵⁴ Auch wenn es sich aufgrund der Komplexität neuronaler Netze bei einer Außenbetrachtung als schwierig darstellt, genau zu erklären, wie und warum ein darauf aufbauender Algorithmus zu einer bestimmten Entscheidung kommt, hat diese nichts mit kreativer Intelligenz zu tun. Seine Entscheidung basiert auf dem Erkennen bestimmter Merkmale, die zusammengefasst als Muster (daher wird unter anderem auch von Mustererkennung [*pattern recognition*] gesprochen) mit einer berechneten Wahrscheinlichkeit einer zuvor erlernten Konstellation entsprechen beziehungsweise ähneln. Der selbstlernende Algorithmus unterscheidet sich von dem nicht selbstlernenden dadurch, dass ihm die zu erkennenden Merkmale nicht vorgegeben wurden. Während also bei einem einfachen Algorithmus der Programmierer oder die Programmiererin definiert, welche Merkmale etwa ein menschliches Gesicht ausmachen (also theoriegeleitet *Regeln* definiert werden), werden selbstlernende Algorithmen trainiert. Man speist sie mit einer Menge von Bildern, bei denen das erwartete Ergebnis bekannt ist – also jeweils definiert ist, ob sie ein Gesicht zeigen oder nicht (es werden also *Ziele* definiert) – daher wird auch von überwachtem Lernen gesprochen.⁵⁵ Der Algorithmus versucht dann, Muster zu identifizieren, die es erlauben, das Bild eines Gesichts von einem Bild ohne Gesicht zu unterscheiden.

-
- 53 Ein Verfahren zum Testen der Stärke von KI ist der sogenannte Turing-Test. Er ist benannt nach dem britischen Mathematiker Alan Turing, für den ein Computer dann als intelligent gelten sollte, wenn dieser schriftlich mit einem Menschen kommunizieren könnte und dieser ihn nicht als Computer, sondern als menschlichen Gesprächspartner identifizieren würde. Diese Anforderung stellt etwa bei dem seit 1991 ausgeschriebenen Loebner Preis und dessen regelmäßigen Wettbewerben die Silbermedaille dar (für Gold kommt die Anforderung gesprochener Sprache hinzu) (vgl. AISB 2019). Bislang ist auch mittelfristig noch kein Gewinn der Silbermedaille absehbar.
- 54 Neuronale Netze bestehen aus künstlichen Neuronen – den Knoten des Netzwerks. Dabei handelt es sich um mathematische Funktionen beziehungsweise Algorithmen. Funktionen verbinden, wie die Synapsen im Gehirn die Neuronen, die Knoten miteinander. Die Netzwerkknoten sind über mehreren Schichten verteilt angeordnet. Jeder Schicht kommt die Aufgabe zu, bestimmte Merkmale oder Muster zu erkennen. Knoten auf der einen Ebene sind mit einem Subset von Knoten auf der nächsten Ebene verbunden. Die Ausgangswerte der unteren Schicht bilden dann die Eingangswerte der nächsthöheren Schicht. Dementsprechend werden die erkannten Muster von Schicht zu Schicht immer komplexer beziehungsweise abstrakter. In der Bilderkennung verlaufen die Abstraktionsgrade so etwa von Helligkeitswerten über Kanten hin zu Linien und Formen, die dann als Objekte interpretiert werden. Mit dem Komplexitätsgrad der Aufgabe steigt dementsprechend die Anzahl an nötigen Schichten, weshalb auch von *deep learning* gesprochen wird (vgl. Spath 2018: 509; Trinkwalder 2016: 131; Ramge 2019b).
- 55 Nicht überwachtes Lernen in Form von *reinforcement learning* [bestärkendem Lernen] wurde 2017 bei AlphaZero beziehungsweise Alpha Go Zero genutzt. Hierbei spielte der Computer einfach ohne menschlichen Input gegen sich selbst, um besser zu werden (vgl. Lovelock 2020: 100; Knight 2019). Eingesetzt wird *reinforcement learning* aber auch in der Robotik, etwa, um von Grund auf selbstständig das Laufen zu erlernen.

den. Dabei geht er nicht intelligent, sondern systematisch vor, indem er die Wahrscheinlichkeit dafür, dass ein Knoten im neuronalen Netz aktiv wird, dessen Parameter oder die Verschaltung zwischen den Knoten so lange anpasst, bis seine Berechnungen dem vordefinierten, erwarteten Ergebnis möglichst nahe kommen.⁵⁶ Es erfolgt somit ein Lernen durch Rückkopplung in Feedback-Schleifen, in denen der Algorithmus systematisch die (initial zufälligen) Parameter des Netzwerks anpasst und neu kalibriert, um die Erkennungsleistung des nächsten Durchgangs zu verbessern.⁵⁷ Es handelt sich um ein iteratives Vorgehen (vgl. u. a. ebd.: 47; Stalder 2017: 177f.).

Wenn die Umsetzung der Entscheidung den digitalen Raum verlässt und, wie die Entscheidungsfindung, ebenfalls technikvermittelt erfolgt, spricht man von autonomen oder automatisierten *cyber-physikalischen* Systemen. Wenn ein autonom fahrendes Auto einem Hindernis ausweicht, folgt es genau dem Prinzip eines Dreischritts aus Sensorik, KI und Aktorik (vgl. auch Spath 2018: 508): »Muster erkennen in Daten [Sensorik]. Erkenntnis durch Statistik und Algorithmen ableiten [KI]. Umsetzung der Erkenntnis in eine Entscheidung durch eine technische Routine [Aktorik]« (Ramge 2019b: 16).

Basis für die Entscheidungen des Algorithmus sind somit das Training mit bekannten oder alten Daten. Eine Gefahr beim Einsatz solcher (selbstlernenden) Algorithmen besteht daher darin, dass sie frühere Fehler fortschreiben beziehungsweise in den Daten vorhandene Diskriminierung replizieren oder sogar verstärken.⁵⁸ Unter anderem aus diesem Grund hat Amazon die algorithmische Unterstützung bei Bewerbungsverfahren wieder eingestellt. Der Algorithmus diskriminierte bei der Bewertung von Bewerbungen diejenigen von Frauen, weil er aus den Daten gelernt hat, dass technikaffine Männer ein größeres Interesse an Amazon als Arbeitgeber hätten. Männer passten besser ins algorithmische Bild des idealen Bewerbers, weil sich in der Vergangenheit häufiger Männer für technische Aufgaben beworben hatten und zugleich auch häufiger eingestellt worden waren (vgl. Holland 2018; Wilke 2018).

Zusammenfassend lässt sich sagen: »Algorithms are not neutral. They are designed by people, with ideologies, biases, and institutional mandates. Algorithms discriminate and make mistakes« (Owen 2015: 200). Daher ist es nicht verwunderlich, dass ethischen Fragen, insbesondere angesichts der Entwicklung Künstlicher Intelligenz, ein zentraler Stellenwert beikommt. Hierzu hat nicht nur die Europäische Kommission im Jahr 2019

56 Daher ist es so schwierig nachzuvollziehen, woran genau ein selbstlernender Algorithmus etwa auf einem Bild ein Gesicht erkennt. »Im Zentrum der Suche steht eine funktionsfähige Lösung, die sich experimentell und in der Praxis bewährt, von der man aber hinterher möglicherweise nicht mehr weiß, warum sie funktioniert, oder ob sie wirklich die bestmögliche Lösung ist« (Stalder 2017: 178). Seyfert (2021: 230) bezeichnet selbstlernende Algorithmen aus diesem Grund als »hochgradig opake Gegenstände« (was aber auch auf andere Algorithmen zutrefte wenn diese etwa proprietär und intransparent seien oder stark von Kontext und Interaktion abhängen).

57 Für eine vertiefende Darstellung der mathematischen Grundlagen neuronaler Netzwerke siehe etwa Trinkwalder (2016).

58 Zu den unterschiedlichen Diskriminierungsrisiken in Verbindung mit der Nutzung von Algorithmen siehe umfassend die Studie von Orwat (2019: xiii), die ebenfalls zu dem Ergebnis kommt: »Viele Diskriminierungsrisiken bei der Entwicklung und Verwendung von Algorithmen resultieren aus der Verwendung von Daten, die frühere Ungleichbehandlungen abbilden.«

ethische Leitlinien einer von ihr einberufenen Expertenkommission vorgestellt (vgl. Europäische Kommission 2019a).⁵⁹ Auch die damalige Bundesregierung hatte mit der Datenethikkommission (vgl. BMI 2018) und der Enquete-Kommission Künstliche Intelligenz – Gesellschaftliche Verantwortung und wirtschaftliche, soziale und ökologische Potenziale (vgl. Deutscher Bundestag 2018a) im selben Zeitraum zwei Gremien beauftragt, sich mit den Herausforderungen durch KI und algorithmische Entscheidungsprozesse kritisch auseinanderzusetzen. Darüber hinaus finden sich erste weiche Formen von Selbstregulierung. So gründete sich im Dezember 2018, ausgehend von einer gemeinsamen Initiative des Bundesverbandes der Personalmanager (BPM) und der Unternehmensberatung hkp///group, der Ethikbeirat HR-Tech. Bis zum Mai 2019 erarbeitete dieser Richtlinien für den verantwortungsvollen Einsatz von KI in der Personalarbeit (vgl. Ethikbeirat HR-Tech 2019).

Bezogen auf die Algorithmisierung werden daher aus gesellschaftspolitischer Perspektive unter den sich neu stellenden ethischen, rechtlichen und sozialen Fragen insbesondere zwei Aspekte intensiv und kontrovers diskutiert: 1. Wie lässt sich Fairness in Algorithmen und algorithmischen Entscheidungsprozessen implementieren und sicherstellen? 2. Inwieweit können, müssen oder sollten ethische Aspekte bei automatisierten Entscheidungen berücksichtigt werden? Hinter diesen Fragen steht unter anderem eine zentrale Gefahr, wie sie etwa Ramge (2019b: 87) hervorhebt: die »Manipulation des Einzelnen« durch Internetkonzerne und Plattformen im »Datenmonopolkapitalismus« auf der einen und den »Missbrauch durch Regierungen« in der »digitalen Diktatur« auf der anderen Seite. Zugleich knüpft sich an Datafizierung und Algorithmizität aber auch die Vorstellung oder zumindest Hoffnung auf positive Potenziale für besseres Regieren.

»Data analytics, in this context, is increasingly viewed and sold as providing a means to more efficiently target and deliver public services and to better understand social problems« (Dencik et al. 2019: 2).

Bezogen auf politische Steuerung werden hierbei zentrale Stellschrauben des Steuerungspotenzials angesprochen. Mehr und bessere Daten sollen nicht nur bei der Problemidentifikation und, davon abgeleitet, der Zieldefinition von Steuerung helfen, sondern das Steuerungswissen per se verbessern. Datenbasiert passgenau auf den Adressatenkreis zugeschnittene Steuerungsinstrumente verheißen eine erhöhte Steuerungsfähigkeit aufseiten des Staates und eine vergrößerte Steuerungswirkung aufseiten der Steuerungssubjekte.⁶⁰ Demgegenüber besteht die Gefahr einer »Herrschaft der Algorithmen«. Weyer (2019: 22) spricht von der Gefahr einer »Algokratie«, Kersting (2018: 87) von der »Algorithmokratie«. Algorithmen können zwar in komplexen, digitalisierten Systemen in der operativen Echtzeitsteuerung notwendig oder hilfreich sein. Dabei geht es jedoch ausschließlich um den normumsetzenden und nicht den

59 Für eine kritische Perspektive auf die Ergebnisse der Kommission, bei denen sich die Interessen der Industrie gegen stärkere zivilgesellschaftliche Vorbehalte durchgesetzt hätten, siehe etwa Köver (2019).

60 Wie daten- und algorithmengetriebene Steuerungsinstrumente eingesetzt werden, wird in Kapitel VI.1.1 näher erläutert.

normsetzenden Teil politischer Steuerung. Schliesky (2020b: 66) weist darüber hinaus darauf hin, dass der Einsatz von Algorithmen zur automatisierten Entscheidung die, in der Demokratie durch Menschen sichergestellte, Legitimationskette unterbricht.

Mit Blick auf politische Steuerung⁶¹ ergibt sich daraus die Dimension der *Governance of Algorithms* (siehe Kapitel V.3.7), also die Frage danach, inwiefern eine Regulierung von Algorithmen notwendig ist (vgl. D'Agostino/Durante 2018). Zum anderen geht es um die Dimension der *Governance by Algorithms* (siehe Kapitel VI.1.1.2), also um die Steuerung durch Algorithmen (vgl. Musiani 2013; Just/Latzer 2017).

II.2.6 Beschleunigung: Die Steigerung der Geschwindigkeit

In Beschreibungen der modernen digitalen Welt mangelt es nicht an Geschwindigkeitsmetaphern: Wir leben in einer *High-Speed Society* (Rosa/Scheuerman 2009a), surfen über den *Information Superhighway*⁶² [Datenautobahn], und Computer handeln Wertpapiere automatisiert im *high-frequency trading* – »Digitalisierung ist Evolution auf Speed«, wie der damalige Vorstandsvorsitzende von Bertelsmann Ostrowski auf seiner Eröffnungs-Keynote zur Kölner DEMXCO feststellte (zitiert nach Winterbauer 2011).

Für die Betrachtung der digitalen Transformationsprozessen ist jedoch die Geschwindigkeit nur eine Seite der Medaille. Deren andere Seite – die Beschleunigung – ist für die Wucht der Veränderung sogar noch bedeutsamer. »Die Beschleunigung von Prozessen und Ereignissen ist ein Grundprinzip der modernen Gesellschaft« (Rosa 2005: 15). Daher können wir »the nature and character of modernity and the logic of its structural and cultural development« nicht verstehen, ohne »the temporal perspective« in unsere Betrachtung miteinzubeziehen (Rosa 2009: 79). Dass die Kategorie der Zeit eine wichtige Rolle spielt, ist auch der Politikwissenschaft nicht fremd. Gerade in der

61 Dafür, dass schon Jahrhunderte vor dem digitalen Zeitalter für Herrscher:innen Daten und Mathematik eine bedeutende Rolle für Machtausübung und -erhalt gespielt haben, siehe etwa Lehner (2018).

62 Der AI Gore zugeschriebene Begriff des Information Superhighway (Gromov 2012), des Information Highway oder schlicht der Infobahn entstand bereits in den 1990er-Jahren in den USA vor dem Hintergrund der steigenden Verbreitung und Nutzung des Internets und der sich ausweitenden digitalen Kommunikation. Interessanterweise fand in Deutschland mit der Datenautobahn ein entsprechend übersetzter Begriff rasch Verbreitung (so stieg etwa die Zahl der Artikel im Spiegel, in denen der Begriff Datenautobahn vorkommt, von 21 seit der Ersterwähnung im Jahr 1993 bis 1994 auf 74 im Zeitraum von 1995 bis 1999 und 63 im Zeitraum von 2000 bis 2004) (vgl. auch Marcuccio 2010). In den anschließenden zehn Jahren sank die Zahl der Artikel dann auf 36 zwischen 2005 und 2009 sowie 43 zwischen 2010 und 2014. Mit der zunehmenden Debatte um den Breitbandausbau ging die Nutzung des Begriffs der Datenautobahn stark zurück. Er fand sich nur noch in 13 Artikeln zwischen 2015 und 2019 und 3 Artikeln in der Zeit von 2020 bis Mai 2022 (Quelle: Spiegel-Archiv; eigene Auszählung unter dem Suchbegriff »Datenautobahn«).

Etwas anders verlief die Verwendung des Begriffs Datenautobahn in Bundestagsreden. Insgesamt fiel der Begriff seit der ersten Erwähnung 1993 in 92 Redebeiträgen. Davon entfielen 49 auf den Zeitraum von 1993 bis 2000. Anschließend sank die Zahl der Beiträge auf 11 im Zeitraum zwischen 2001 und 2010, um dann mit dem Einsetzen der Breitbanddiskussion wieder auf 31 in den Jahre 2011 bis 2020 anzusteigen (Quelle: OpenDiscourse.de; eigene Auszählung unter dem Suchbegriff »Datenautobahn«).