

Guido Arnold | Margarete Jäger | Helmut Kellershohn (Hg.)

(Post-)Pandemische Normalitäten

Zu den gesellschaftlichen Auswirkungen
der Corona-Krise



Edition des Duisburger Instituts für Sprach- und
Sozialforschung im *UNRAST* Verlag, Münster



Korrelation versus Kausalität

Chris Anderson (ehemaliger Chefredakteur des Technologie-Magazins *Wired*) läutete bereits 2013 in einem Essay »Das Ende der Theorie« ein. Er schrieb, man brauche keine semantische oder kausale Analyse mehr – eine statistische reiche völlig aus:

»Wir leben in einer Welt, in der riesige Mengen von Daten und angewandte Mathematik alle anderen Werkzeuge ersetzen, die man sonst noch so anwenden könnte. Ob in der Linguistik oder in der Soziologie: Raus mit all den Theorien des menschlichen Verhaltens! Vergessen sie Taxonomien, die Ontologie und die Psychologie! Wer weiß schon, warum Menschen sich so verhalten, wie sie sich gerade verhalten? Der springende Punkt ist, dass sie sich so verhalten und dass wir ihr Verhalten mit einer nie gekannten Genauigkeit nachverfolgen und messen können. Hat man erst einmal genug Daten, sprechen die Zahlen für sich selbst.«⁵

Mehr noch: Anderson diskreditierte darin *Nicht-BigData*-Ansätze wie etwa die klassische Theoriebildung in der Physik als »*Schöne Geschichten-Phase einer Disziplin, die an Datenhungerr litt*«. ⁶ Beinahe könnte der Eindruck entstehen, Anderson habe recht: Die erfolgreichste Variante künstlich-intelligenter Sprachübersetzungs-Software von Google verzichtet fast vollständig auf grammatikalische Vorgaben der beteiligten Sprachen. Die Übersetzung basiert ausschließlich auf ausreichend vielen Textdaten, die in einen selbst-lernenden Algorithmus zur Sprachmustererkennung eingespeist werden. Google versucht also gar nicht erst, eine Sprache >zu verstehen< – das heißt, per Abstraktion Sprachregeln abzuleiten und sich somit eine Theorie der Sprache zu erarbeiten. Stattdessen wird die Wahrscheinlichkeit einer treffenden Übersetzung ganzer Wortgruppen durch zwei Dinge erhöht: die stetige Erweiterung der Vergleichsdatenbanken bereits getätigter Übersetzungen und eine darauf basierende, sich automatisch anpassende Neugewichtung >erkannter< Übersetzungsmuster. Nach Milliarden dementsprechend >übersetzter< Texte könnten KI-Enthusiast*innen behaupten, diese erlernten Gewichte (Korrelationen) gäben eine Art >Grammatik< der Sprachübersetzung wieder.

Der entscheidende Punkt ist: Diese >Grammatik< ist nicht extrahierbar – die Allgemeinheit profitiert nicht in einer von der künstlich-intelligenten Maschinerie abstrahierbaren Form von diesem Übersetzungswissen. Das Wissen ist nicht von der immer größer werdenden Sprachmusterdatenbank abzutrennen: Wir lernen nichts über das Wesen der beteiligten Sprachen. Googles Übersetzungsassistent ist eine >Black Box<. Google häuft Herrschaftswissen an und behauptet, einen Beitrag zur allgemeinen Verständigung und zur Wissensvermehrung zu liefern – für alle frei zugänglich.

⁵ Chris Anderson: Das Ende der Theorie. Die Datenschwemme macht wissenschaftliche Methoden obsolet, in: Geiselberger, Heinrich/ Moorstedt, Tobias (Hg.): Big Data – Das neue Versprechen der Allwissenheit, Berlin 2013, 124-131.

⁶ Ebd., 127

Joseph Weizenbaum, selbst KI-Pionier in den 1960er Jahren und später einer ihrer prominentesten Kritiker, warf den KI-Technokrat*innen die bewusste Verschleierung des Unterschieds zwischen (phänomenologischer) *Beschreibung* und einer *Theorie*, die semantische bzw. kausale Zusammenhänge aufzeigt, vor. Besonders fatal werde dies bei der Modellierung des Menschen mit seinen Persönlichkeitseigenschaften. Hier geben KI-Enthusiast*innen das Datenabbild eines Menschen für den Menschen selbst aus. Das meint, die Beschränkungen einer modellhaften Beschreibung nicht nur zu vergessen, sondern bewusst zu vertuschen. Weizenbaum prognostizierte fatale Folgen bis hin zum Verlust der Eigenständigkeit in einer von Künstlicher Intelligenz fremdbestimmten Welt digitaler Dienste.

Die Kausalität ist, anders als die Korrelation, ihrem Wesen nach überprüf-, hinterfrag- und angreifbar. Sie ist nicht proprietär und eignet sich daher weniger zum Aufbau von Machtgefälle.

Diskriminierende Algorithmen

KI-basierte Analyse-Software wird mittlerweile standardmäßig zur Verhaltenserschaffung und -lenkung eingesetzt. In vielen Lebensprozessen begegnen wir mehr oder weniger aufwändig programmierten, künstlich intelligenten Assistenten – in jeder Spracherkennungssoftware beispielsweise oder bei der individuellen Profilerstellung durch Datensammeldienste wie *Google*, *facebook*, *Palantir* oder *Amplitude*. Wer erhält welche individuellen Bonusmöglichkeiten beim (Online-)Einkauf, bei der Kranken- oder Kfz-Versicherung?

Doch dabei bleibt es nicht. Zunehmend regelt die algorithmische Verhaltensanalyse auch die Beziehung des Einzelnen zum Staat – zu Bildungs- und Gesundheitseinrichtungen, zum Sicherheitsapparat und zur Politik.⁷ Auf Grundlage von prognostizierten Verhaltensweisen, Risiken oder Effizienzpotenzialen wird Bevölkerung automatisiert *ungleich behandelt*.

Die künstlich intelligente Vorhersage der interessierenden (sensiblen) *Verhaltensdaten* erfordert dazu eine möglichst detaillierte Erfassung und anschließende Bewertung von sogenannten *Hilfsdaten*.

So werden etwa aus den *Hilfsdaten* eines individuellen Browserverlaufs die sensiblen *Verhaltensdaten* zur sexuellen Orientierung einer Person mit hoher Trefferquote vorhergesagt. Genau wie bei Amazons künstlich intelligenter Analyse von Bewerbungsmappen benötigt die selbstlernende KI *Trainingsdaten* – also Personen, die sowohl ihre Browserverläufe als auch Angaben zu ihrer sexuellen Orientierung offenlegen. Anhand dieser Trainingsdaten lassen sich Muster erkennen – also

⁷ Karen Hao: AI is sending people to jail—and getting it wrong, MIT Technology Reviews; <https://www.technologyreview.com/2019/01/21/137783/algorithms-criminal-justice-ai> (21.09.2019).

statistische Häufigkeiten (Korrelationen), die angeben, wie oft welche Eigenheiten des Browserverlaufs mit welchen sexuellen Vorlieben verknüpft sind.

Diese Korrelationen geben jedoch nur Auskunft über die statistische Gruppe dieser Trainingspersonen. Eine induktive Verallgemeinerung über die Trainingsdaten hinaus ist die (unzulässige) Grundlage der prädiktiven Verhaltensanalyse: Die ermittelten Korrelationen werden genutzt um aus beliebigen (anderen) Browserverläufen die sexuelle Vorlieben der zugehörigen Internetnutzer*innen vorherzusagen.

Es ist klar, dass 1.) diese Ähnlichkeitsgruppierungen nach dem *people-like-you* Prinzip keine kausalen Verknüpfungen zwischen Browserverlauf und sexueller Orientierung zulassen. Neben der 2.) *unzulässigen Verallgemeinerung* dieser statistischen Gruppeneigenschaft über die Gruppe der Trainingspersonen hinaus findet mit der nun abgeleiteten Verhaltensvorhersage jedoch auch eine 3.) *unzulässige Vereindeutigung* auf eine Einzelperson statt. Ohne jede Unschärfe ordnet die prädiktive Analyse einer Einzelperson exakt eine Schublade der sexuellen Orientierung zu. Letzteres wird bei der Kritik an der prädiktiven Verhaltensanalyse oft vergessen und lediglich der sogenannte *Bias* der Trainingsdaten (Verzerrung der statistischen Gewichte durch nicht-repräsentative Auswahl von Trainingspersonen) bemängelt.

4.) Das eigenständige Entwickeln von stetig neu ausgerichteten Ähnlichkeitsmustern (durch die selbstlernende KI) bringt es mit sich, dass die Unterscheidung von Personen(-gruppen) dynamisch und nicht statisch erfolgt: Erfassung und Bewertung erfolgen nicht mehr nach fest definierten Kategorien – der Schubladenschrank ist in Ausprägung und Anzahl der >Schubladen< variabel. Eine Schufa, die eine Einstufung der Kreditwürdigkeit mit feststehender Gewichtung einmal benannter Kriterien erlaubt, wäre ein hoffnungslos veraltetes finanzpolitisches Instrument.

Über prädiktive Verhaltensanalyse haben Mediziner*innen der University of Pennsylvania aus Anzahl und Inhalt von Facebook-Beiträgen (den Hilfsdaten) die Neigung zu Depressionen, Psychosen, Diabetis, Bluthochdruck (als sensible Verhaltensdaten) vorhergesagt.⁸ Facebook selbst geht in den USA so weit, aus Abweichungen beim individuellen Schreibverhalten eine etwaige Suizidgefahr zu detektieren und in Reaktion darauf nicht nur regionale Hilfsangebote einzublenden, sondern in besonders akuten Fällen auch Sicherheitsbehörden zu kontaktieren, die sich ggfs. Zutritt zur Wohnung der Nutzer*in verschaffen.⁹

8 R. M. Merchant et al.: Evaluating the Predictability of Medical Conditions from Social Media Posts, in: PLOS ONE 14 (2019), Nr. 6. DOI: 10.1371/journal.pone.0215476.

9 Dies Funktionalität ist in den USA seit 2017 aktiviert worden. In der EU ist sie aus datenschutzrechtlichen Gründen deaktiviert.

Programmatische Ungleichbehandlung

KI-basierte Prognosesysteme stellen ein *sozio-technisches* Instrument (z.B. für das Bevölkerungsmanagement) dar, da sie zum einen die Trainingspersonen involvieren und zum anderen auf der Akzeptanz derer basieren, die zwar glauben, ihre sensiblen Persönlichkeitsdaten gut zu schützen, aber ihre vermeintlich bedeutungslosen Hilfsdaten freizügig im Netz hinterlassen und damit die Datenbasis für die Prognose bieten. In diesem Prozess ist nicht jede*r ihre/seine eigene abgeschlossene Datenbasis und nur für seine eigene Vorhersagbarkeit per KI und die daraus mögliche Verletzung der Privatsphäre verantwortlich. Die Diskriminierung (Unterscheidbarkeit) einzelner ist erst über die kollektive Datensammlung durch viele möglich.¹⁰ Damit ist diese Diskriminierung auch kollektiv zu verantworten.

*>Normale< Nutzer*innen, die >nichts zu verbergen haben< ermöglichen mit ihren Daten prädiktive Algorithmen, die andere Bevölkerungsgruppen als Abweichler*innen markieren.*

Prädiktive Analysen stabilisieren die prognostizierte >Realität<. Durch ihre Vereindeutigung in der (eigentlich statistisch unscharfen) Vorhersage schreiben sie selbstverstärkend (per Feedback-Schleife) das Vorhergesagte fest. Die durch KI-Verhaltensvorhersage ermöglichte Ungleichbehandlung verschiedener Ähnlichkeitsgruppen zementiert eine Unterteilung in *unsichtbare* soziale Klassen. Unsichtbar deshalb, da die Unterscheidungskriterien keine demografisch nachvollziehbaren, sondern per KI-Optimierung vollständig intransparent sind.

R. Mühlhoff liefert eine biopolitische Interpretation des damit möglichen algorithmischen Bevölkerungsmanagements: »Jedes Individuum wird individuell erfasst, angesprochen und behandelt, dabei aber doch nicht aus dem Glaskäfig einer virtuellen Vergleichsgruppe entlassen.«¹¹ In der Konsequenz dieser computerisierten Ungleichbehandlung konstatiert er eine Entsolidarisierung hin zu einer Ethik der sozio-ökonomischen Selektion – jede*r wird behandelt, wie sie/er es verdient. Wer ungesünder lebt, zahlt zukünftig mehr für die Krankenversicherung. Die alltägliche Gewöhnung an eine >passgenaue< Ungleichheit etabliert sich schleichend als neues Gerechtigkeitsempfinden. Die Transformation von einer Ungleichbehandlung als Ungerechtigkeit zur neuen Gerechtigkeit vollzieht sich passend zum Solutionismus der angewandten Methode ganz unideologisch.

10 Der ehemalige NSA-Chef Keith Alexander formulierte dies kontraintuitiv aber korrekt folgendermaßen: »in order to find a needle in the haystack, you have to collect the haystack.«

11 Rainer Mühlhoff: Prädiktive Privatheit: Warum wir alle »etwas zu verbergen haben«, in: #VerantwortungKI – Künstliche Intelligenz und gesellschaftliche Folgen, hrsg. von Christoph Marksches und Isabella Hermann, Bd. 3/2020, Berlin-Brandenburgische Akademie der Wissenschaften.

Wer hat Zugang zu welchem Gebäude oder Stadtteil? Dies *generell* per KI zu regulieren, erscheint uns heute (noch!) als eine dystopische Übertreibung.¹² Die in China umgesetzte Idee einer App-gesteuerten Lockerung des allgemein verfügbaren Lockdowns in der Coronakrise macht jedoch genau das: Die virologische Gefährdung, die von einem Individuum ausgeht, wird anhand von individueller *Mobilität*, *Anzahl an Kontakten* (beides quantifizierbar über ein Smartphone-Tracking bzw. -Tracing) und digitaler Einstufung des >Immunitätsgrades< >bemessen<. Die intransparente Bewertung mündet dann in drei unterschiedliche Bewegungsfreiheitsgrade: (grün=) freier Zugang z.B. zum Einkaufszentrum, (gelb=) kein Zugang, (rot=) kein Zugang mit sofortigen Quarantäneauflagen.

Ist es in ordnungspolitischer Denkweise nicht konsequent, den Menschen generell als multidimensionalen Risikofaktor wahrzunehmen und damit die Kategorie der *Gefährder*in* entlang unterschiedlicher Risiken weiter zu differenzieren und gemäß >Gefährder*innengrad< unterschiedliche Einschränkungen zu verordnen? Bei sehr vielen Gefährdungsstufen landen wir beim *Scoring*. Denn es ist für die relationale Unterscheidung unerheblich, ob mensch hochdifferenzierte Einschränkungen per Malus verhängt oder in der Umkehrung soziale Teilhabe per Bonus diversifiziert. Die Intransparenz der Einstufung ermöglicht so im genannten Beispiel der Quarantäneampel den missbräuchlichen Ausschluss missliebiger Personen aus weiten Teilen des öffentlichen Lebens – keine dystopische Projektion in die Zukunft, sondern >smarte< Realität im Jahr 2022.¹³

Kategorisierung ohne Kategorien

Das *Scoring* erlaubt eine KI-basierte Diskriminierung nicht mehr nur entlang vermeintlicher Gruppenmerkmale, sondern bis zur Ebene einzeln unterscheidbarer Individuen. *Soziale Scoring-Systeme* werden auch außerhalb von China immer populärer. Verschiedene Wohn-, Job-, Kredit- oder Mobilitätsangebote gelten nur für Teilnehmer*innen mit genügend hohem >Score< (= erworbene Punkte durch belohnenswertes Verhalten).

Es entsteht dann kein grobes Schubladensystem einzelner Klassen, sondern eine feinstkörnige Individualisierung. Feinstkörnig in dem Sinn, dass die Differenzierung entlang so vieler Parameter durchgeführt wird, dass sie *vollständig* ist: Eine weitere Unterscheidung über noch mehr Parameter würde das Schubladensystem

12 Für einige Orte ist dies schon Realität und das nicht nur in China. Zum Beispiel, wenn lokale US-Polizeibehörden ihre Datenbanken vernetzen und Personengruppen aus bestimmten Gebieten fernhalten oder sogenannte *gated communities* sich komplett abschotten und Zutritt nur per Amazon Ring [einer >intelligenten< Türklingel mit netzgekoppelter Kamera] gestatten.

13 Vgl. <https://www.tagesschau.de/ausland/asien/china-app-corona-kontrolle-101.html>

nicht weiter verfeinern, sondern es lediglich noch unwahrscheinlicher machen, dass zwei Menschen in derselben ›Kategorie‹ landen, also exakt gleiche Teilhabemöglichkeiten (= gleiche Punktezahl) zugewiesen bekommen.

Der Begriff ›Kategorie‹ ergibt dann keinen Sinn mehr. Auch der Begriff der ›Klasse‹ wird deformiert. Klassenzugehörigkeit lässt sich in einer per Score und Ranking verfassten Gesellschaft nicht mehr sinnvoll durch eine eindimensionale Bewertung von (lohn-abhängiger) Arbeit definieren. Stattdessen führt die Erfassung von zigtausenden Mustern von Verhalten, Kontakten, Einstellungen und Wünschen auch jenseits der Lohnarbeit zu einem niedrigen Score – das ist die neue vermeintlich ›klassenlose‹ Deklassierung und sie ist in vielen Metropolen Chinas bereits lebensbestimmende Realität.

Die ›Individuierung‹ per Score wird von Solutionist*innen vorangetrieben und findet entlang kapitalistisch motivierter Bewertungskriterien statt. Sie ordnet aber nicht nur den Markt neu, sondern auch die politische Administration. Wir werden daran gewöhnt, dass regelnde Verordnungen *nicht mehr für alle gleich* sind. Das bisherige Steuersystem z.B. ist so verfasst, dass sich Steuerklassen und der letztendliche Steuerbetrag einer Person über einen simplen klassischen Algorithmus berechnen lassen – also eine einfache Abfolge von statischen *Wenn-Dann*-Beziehungen. In einer feinkörnigen Scoring-Gesellschaft hingegen gelten für jede*n andere Regeln. Es gibt keine Steuerklassen. Stattdessen erhält jede*r individuelle Verhaltensempfehlungen von einem KI-basierten digitalen Assistenten. Diese Handlungsempfehlungen umzusetzen oder zu ignorieren wird mit Bonus- oder Malus-Punkten belohnt oder sanktioniert. Das ergibt maximale Individuierung in dem Sinne, dass niemandes Situation mit der eines anderen vergleichbar ist. Sogar die jeweiligen Bedingungen, gemäß derer jede*r einzelne Punkte sammelt, sind damit nicht vergleichbar. Das befördert maximale Entsolidarisierung. Ein kollektives Aufbegehren wird erschwert. Gruppen von (vergleichbar) Betroffenen lassen sich nur aufwändig bilden. Wer will Ungleichbehandlung beklagen oder gar skandalisieren, wenn sie offen konstitutiv für das System ist?

Intransparenz als Basis für Selbstoptimierung

Ein wesentliches Merkmal KI-basierter Systeme zur Verhaltenslenkung ist die *Undurchsichtigkeit der Erfassungs- und Bewertungskriterien*. Die Kategorien, gemäß derer Verhalten erfasst werden, sowie ihr Einfluss auf die Gesamtbewertung bleiben bewusst (z.B. Schufa-)Geheimnis. Mehr noch: Die Muster, nach denen sich besonders effizient Verhalten unterscheiden lässt, verändern sich im Rahmen einer selbstlernenden KI. Je mehr Daten ins System eingespeist werden, desto ›treffen-der‹ findet die KI eine Unterscheidung (gemäß ihrer Lernvorgaben) *wesentlicher* Verhaltensmerkmale. Die zu bewertenden Individuen können diese *dynamische*

Kategorisierung gar nicht kennen. Sie können lediglich errahnen und spielerisch (Gamification) erforschen, welches Verhalten ihre (momentane!) Punktezahl wie stark beeinflusst. Eine optimale Voraussetzung dafür, sich in Unkenntnis der Bewertungsmodalitäten und in der Hoffnung auf einen besseren Score umfassend selbst zu optimieren. Denn das Scoring-System ist über ein (nicht einsehbares) Ranking konkurrenzbasiert.

Die Situation ist noch komplexer: Selbst die Informatiker*innen, die das Bewertungsprogramm entwickelt haben, können nur zu Beginn eine Aussage darüber treffen, welches Verhalten zu welchem Score in der KI-Bewertung führt. Nach (selbstlernender) Weiterentwicklung des Programms verändern sich die relevanten Muster und ihre Gewichtung für den Score. Die Informatiker*in kennt dann die >Gewichte< ihrer eigenen Individuierungssoftware nicht mehr. Diese Unkenntnis ist nicht Ausdruck ihrer Unfähigkeit, sondern systemisch bedingt und eher als ein Maß für die Effektivität einer KI-basierten Optimierung ohne starke Vorgaben zu verstehen.

Diese Freiheit der KI ist Vorzug und Makel zugleich. Selbstlernende neuronale Netze arbeiten besonders gut, wenn sie ihre Optimierungsmuster so >frei< wie möglich selbst entwickeln können. Die Details künstlich intelligent getroffener Entscheidungen entziehen sich jedoch einer menschlichen Nachvollziehbarkeit und damit auch ihrer Vermittelbarkeit.

Gesellschaftliche Gerechtigkeitsvorstellungen basieren auf Vergleichbarkeit und suchen Ungleichheiten perspektivisch abzuschaffen. Was bedeutet es für eine Gesellschaft, wenn sie es Technokratie und Solutionismus überlässt, Regeln zu entwerfen, die auf Ungleichbehandlung setzen und weder beständig noch vermittelbar, ja, nicht einmal bekannt sind?

›Verantwortungsvolle KI – ist eine progressive Wendung möglich?

Könnten die anfänglichen Optimierungsparameter einer KI nicht so gewählt werden, dass sich eine derartige digitale Assistenz an einer >Mehrung des Gemeinwohls< orientiert? Die Antwort muss nein lauten. Künstliche Intelligenz steigert nicht das Gemeinwohl eines komplexen sozialen Gefüges. Denn die KI ist schlicht *begriffslos* bezüglich der *Bedeutung* eines Gemeinnsinns. Sie hat keine Idee davon, was das *Wesen* einer Gemeinschaft sein könnte. Egal wie viele Parameter ihre Optimierungsfunktion aufweist, egal wie gut und umfangreich die ihr zur Verfügung gestellten Trainingsdaten auch sein mögen. Die KI befördert vielmehr den Zerfall der Öffentlichkeit.¹⁴ Insbesondere die ihr immanente *Ungleichbehandlung* unterminiert

14 Zyenep Tufekci: Engineering the public – big data, surveillance and computational politics, 02.07.2014, <http://doi.org/10.5210/fm.v19i7.4901>

eine politische Öffentlichkeit. Die bequeme Bevormundung digitaler Assistenz zieht eine effektive Einschränkung der Willensfreiheit¹⁵ nach sich, die Hannah Arendt als konstitutiv für politisches Handeln nennt: Die Fähigkeit zu reflektieren und aktiv im Sinne eines Gemeinwohls zu partizipieren, um darüber neue Formen des Gemeinsinns zu imaginieren. Antoinette Rouvroy sieht in einer algorithmischen Gouvernementalität gar den Tod der Politik.¹⁶

Eine Reduktion dieses intrinsischen Problems der künstlichen Intelligenz auf die reine *Anwendung* dieser Technologie unterschätzt die Tragweite und unterschlägt die Dynamik der technologischen Innovation als (Herrschafts- und Subjektivierungs-) Prozess. Darunter leidet derzeit auch ein Regulierungsversuch der EU unter dem Namen *Artificial Intelligence Act* (AIA)¹⁷ mit dem Ziel, Rahmenbedingungen für eine »vertrauensvolle KI« zu schaffen. Laut Gesetzentwurf richtet sich die Notwendigkeit zur Regulierung einer KI-Software daran, ob dieser ein »nicht akzeptables« oder ein »hohes Risiko« zugeschrieben wird. Biometrische Identifizierung per KI fällt in die letztgenannte Kategorie und ist weiterhin erlaubt. Die Betreiber müssen lediglich eine hohe Datenqualität und ein Risikomanagement nachweisen, haben Transparenzpflichten gegenüber Nutzenden und müssen eine technische Dokumentation vorlegen. Soziale-Kredit-Systeme (wie wir sie aus China kennen) gelten vorgeblich als »nicht akzeptabel« und sollen verboten werden. Dies darf in dieser Konsequenz bezweifelt werden, denn in einzelnen Branchen wie dem Versicherungswesen sind sie modular bereits seit über drei Jahren im Einsatz – ohne Beanstandung. Die oben diskutierte Emotionserkennung gilt als Anwendungsgebiet mit »begrenztem Risiko«. Hier müssen die Nutzerinnen lediglich darauf hingewiesen werden, dass solche automatisierten Systeme zur Anwendung kommen. Es ist davon auszugehen, dass der Regulierungsversuch ins Leere laufen wird, da die zu kontrollierenden Unternehmen selbst die Abschätzung der gesellschaftlichen Folgen vornehmen sollen und die EU großes Interesse bekundet, die Entwicklung von KI-Systemen mit dem AIA nicht behindern zu wollen. Was bleibt ist eine Art EU-Gütesiegel, von dem sich die EU einen Wettbewerbsvorteil erhofft – im geradezu erdrückenden Konkurrenzkampf mit den deutlich investitionsstärkeren USA und China. Ein Gütesiegel, welches sich plakativ von Chinas Ziel einer quasi-totalitären, digitalen Vermessung und Bewertung aller Individuen in Echtzeit abgrenzt, ohne den hoch effektiven Methoden eines KI-basierten, biopolitischen Bevölkerungsmangements abzuschwören.

15 Guido Arnold: Nudging – die politische Dimension psychotechnologischer Assistenz, in: DISS-Journal 43/2022.

16 <https://www.greeneuropeanjournal.eu/algorithmic-governmentality-and-the-death-of-politics/>

17 https://ec.europa.eu/info/sites/info/files/commission-white-paper-artificial-intelligence-feb2020_en.pdf