

NETZWERK
WISSENSCHAFT



Bit für Bit in die Zukunft

Künstliche Intelligenz
in Wissenschaft und
Forschung

Angela Borgwardt



FRIEDRICH
EBERT 
STIFTUNG

Hinweis:

Für den optimalen Onlinegebrauch wurde diese Version der Publikation mit **Hyperlinks** ausgestattet.

Sämtliche im Text vorkommenden URLs sind direkt verlinkt. **Sie sind entsprechend gekennzeichnet.**

Bit für Bit in die Zukunft

—

Künstliche Intelligenz
in Wissenschaft und
Forschung

Angela Borgwardt

**Schriftenreihe des
Netzwerk Wissenschaft**

INHALT

Handlungsempfehlungen	05
Einführung	14
Künstliche Intelligenz: große Chancen, große Risiken	25
Beispiele aus der Anwendung	39
— KI in der Industrie	39
— KI in der Medizin	54
— KI als Kommunikationspartner	62
Die Rolle von KI in der Wissenschaft	72
Politik und KI – eine geregelte Beziehung?	82
Impressum	91

HANDLUNGSEMPFEHLUNGEN

1. Verantwortungsvollen Umgang mit KI erreichen

In Deutschland und Europa sollte ein Konzept der KI (Künstliche Intelligenz) umgesetzt werden, das auf ein demokratisches, offenes und vielfältiges Gemeinwesen ausgerichtet ist – in Abgrenzung zu Konzepten in den USA oder in China, wo KI-Forschungsprogramme sich in der Regel an sicherheitspolitischen bzw. militärischen Fragen orientieren oder in ein autoritäres Gesellschaftskonzept eingebettet sind.

Eine KI „Made in Europe“ sollte sich am Wohl der Gesellschaft orientieren und auf eine Entlastung des Menschen, eine Verringerung sozialer Ungleichheiten und eine nachhaltige Entwicklung ausgerichtet sein. Dazu gehört, Menschen nicht als Datenlieferant_innen oder manipulierbare Käufer_innen zu betrachten. Vielmehr sollte ein ganzheitliches Menschenbild verfolgt werden, das den Menschen in allen gesellschaftlichen Bereichen in den Mittelpunkt stellt: Es bedarf aufgeklärter Bürgerinnen und Bürger, die über Kompetenzen verfügen, KI-Systeme zu verstehen, zu nutzen und mit ihnen zu interagieren. Übergreifendes Ziel sollte es deshalb sein, Menschen über KI (kognitiv) zu befähigen und in der Gesellschaft einen verantwortungsvollen Umgang mit KI zu etablieren, der sich an demokratischen Werten, Menschenrechten und dem Prinzip der Rechtsstaatlichkeit orientiert.

2. Erklärbarkeit und Verständlichkeit von KI anstreben

Eine wichtige Grundlage für einen verantwortungsvollen Umgang mit KI ist Erklärbarkeit: Die Verarbeitungsschritte eines KI-Systems müssen verständlich und nachvollziehbar sein, um die Gründe zu verstehen, die zu einer Entscheidung geführt haben. Nur dann kann der Mensch daraus sinnvolle Schlüsse ziehen und das System weiter optimieren, und nur dann ist es möglich, Vertrauen und Sicherheit zu schaffen, Gerechtigkeit und Transparenz zu verwirklichen sowie ethische und rechtliche Unabhängigkeit zu wahren. Eine erklärbare KI ist für Entwickler_innen und Anwender_innen wichtig, aber auch für Nutzer_innen und die von KI-Systemen Betroffenen. Sie müssen den Einsatz und die Auswirkungen von Algorithmen verstehen, um informiert handeln zu können und systematische Verzerrungen zu erkennen, z.B. Diskriminierungen, die aus automatisierten Entscheidungen resultieren. Bei einem gemeinwohlorientierten Einsatz und der gesellschaftlichen Akzeptanz von KI spielt die Aufklärung der Nutzer_innen eine zentrale Rolle. Deshalb muss intensiv auf ein allgemeines Verständnis von KI-Systemen hingewirkt werden.

Der Mensch gibt einem KI-System durch das Programmieren eine Richtung vor und die Maschine kommt über Algorithmen und Kombinatorik zu bestimmten Schlüssen. Dadurch werden Menschen in die Lage versetzt, bessere Entscheidungen zu treffen. Auch die Wissenschaft kann durch KI gewinnen, z.B. indem Forscher_innen KI-Verfahren dazu einsetzen können, den Wissensstand oder die Gesamtheit von Publikationen in einem Fachgebiet systematisch zu erschließen und zu ordnen.

In komplexen Entscheidungssituationen können KI-Systeme zwar neue Phänomene entdecken und dem Menschen Vorschläge machen, doch sollte ein Mensch diese automatisch generierten Schlussfolgerungen aufnehmen, reflektieren, bewerten und letztlich die Entscheidung treffen.

3. KI-Forschung breit und divers fördern

Um die KI-Wissenschaft in Deutschland zu stärken, müssen die staatlichen Investitionen in KI-Forschung an den Hochschulen erhöht werden. Dabei sollte möglichst divers gefördert werden, d.h. sowohl anwendungsorientierte Forschung als auch Grundlagenforschung ohne direktes Ziel der Verwertbarkeit. Auch sollte sich der Fokus nicht nur auf maschinelles Lernen/Deep Learning richten, sondern gezielt auch auf andere Bereiche gegen den Trend, z.B. Wissensrepräsentation. Zudem gilt es, KI-Forschung auf weitere gesellschaftliche Anwendungsbereiche auszuweiten, indem Schwerpunkte wie Mobilität und Gesundheit durch derzeit vernachlässigte Bereiche wie Bildung ergänzt werden.

Wichtig ist, dass KI-Forschung nicht nur an wenigen Exzellenzstandorten stattfindet, sondern das ganze Wissenschaftssystem umfasst. Entsprechend sollten vielfältige Disziplinen, Themenbereiche und Wissenschaftseinrichtungen gefördert werden, um unterschiedliche Perspektiven auf das Thema zu berücksichtigen. Dies ist eine wichtige Voraussetzung, um die mit KI verbundenen Fragestellungen und Herausforderungen angemessen zu bearbeiten.

Um eine humanistische und gemeinwohlorientierte KI zu gestalten, muss ethische und gesellschaftliche Forschung einen deutlichen höheren Stellenwert erhalten als bisher. Es bedarf einer stärkeren kritischen Reflexion von KI, die sich auch mit möglichen Gefahren wie z.B. Datenmissbrauch und Fragen der Verantwortung beschäftigt. Notwendig ist deshalb eine verstärkte Förderung der Geistes- und Sozialwissenschaften, einschließlich der Digital Humanities. Dabei sollten diese nicht nur als dialogische Forschung gefördert werden, also im Sinne einer Vermittlung und Akzeptanzförderung von KI, sondern auch im Bereich der KI-Forschung selbst. Bei der Umsetzung der KI-Strategie der Bundesregierung sollte somit darauf geachtet werden, in diesem Bereich Forschung zu fördern und auch explizit Professuren in ethischer und gesellschaftlicher Forschung einzurichten.

4. KI-Expertise im gesamten Wissenschaftssystem aufbauen

Im Rahmen ihrer KI-Strategie will die Bundesregierung hundert zusätzliche KI-Professuren einrichten, um die KI-Forschung zu stärken. Die Besetzung der Professuren sollte mit Augenmaß stattfinden, da es ausreichend Zeit braucht, um qualifizierte KI-Wissenschaftler_innen in dieser Zahl zu finden. Deshalb sollten die Hochschulen die Möglichkeit haben, diese Stellen sukzessive in den nächsten Jahren zu besetzen. Entscheidend ist dabei, klare Qualitätskriterien für die KI-Professuren zu benennen und die Berufungen konsequent daran auszurichten. Dabei muss auch auf Diversität geachtet werden, d.h. es sollten ausreichend weibliche und junge Wissenschaftler_innen zum Zug kommen.

Die Schaffung von KI-Professuren muss durch weitere Aktivitäten auf verschiedenen Ebenen ergänzt werden, da nachhaltige KI-Expertise im Wissenschaftssystem nur „von unten“ aufgebaut werden kann. Deshalb ist es notwendig, zügig Maßnahmen in der Nachwuchsförderung und im Ausbildungsbereich umzusetzen. Zum einen sollten Doktorand_innen- und Postdoc-Programme in der KI-Forschung aufgelegt werden, zum anderen sollten an den Schulen verstärkt digitale Schlüsselkompetenzen vermittelt und die Schüler_innen mit Informatik und KI vertraut gemacht werden.

5. KI rechtlich und politisch regulieren

Um Fehlentwicklungen und Missbrauch von KI entgegenzuwirken, bedarf es geeigneter rechtlicher Rahmenbedingungen. Diese müssen sicherstellen, dass sowohl Rechtssicherheit als auch ausreichend Spielraum für Innovationen gegeben ist.

Eine wichtige Frage bei KI-Systemen ist, wer bei Entscheidungen die Verantwortung trägt, insbesondere bei Fehlern. Anwender_innen oder Nutzer_innen müssen wissen, wer praktisch zur Verantwortung gezogen werden kann. Der Gesetzgeber sollte hier für Haftungsklarheit sorgen und deshalb eine Person bzw. eine Stelle festlegen, an die sich der bzw. die Einzelne wenden kann. Teilweise kann das bestehende Recht bei KI angewendet werden, doch ist es sinnvoll, bei grundlegenden Innovationen wie z.B. autonom fahrenden Autos neue gesetzliche Regelungen zu schaffen, in denen ein konkreter Anspruch formuliert und speziell bestimmt wird, wer die Verantwortung tragen soll. Das würde den Bürger_innen mehr rechtliche Klarheit geben.

Kooperationen von großen Digitalunternehmen und Hochschulen nehmen im KI-Bereich immer mehr zu. Diesen Bereich muss die Politik stark regulieren, damit die Unabhängigkeit der Forschung gewahrt bleibt, die Zusammenarbeit auf Augenhöhe und nach nachvollziehbaren Regeln stattfindet und die Gewinne nicht ungleich verteilt werden bzw. die Zusammenarbeit zulasten der Hochschulen geht. Es müssen klare Regelungen für Kooperationen zwischen Hochschulen und Wirtschaft getroffen werden, die für die Öffentlichkeit auch transparent sind.

6. Neue Datenökonomie mit Datensouveränität verbinden

Eine wesentliche Frage ist, wie zukünftig die Verknüpfung und Verfügbarkeit von Daten wirtschaftlich organisiert werden soll. Im Bereich der Datenökonomie ist derzeit auf globaler Ebene eine Monopolisierung von riesigen, personenbezogenen Datenmengen auf großen Plattformen festzustellen: Die digitalen Märkte werden inzwischen von wenigen internationalen Datenkonzernen wie z.B. Google und Facebook dominiert und die Menschen werden nur noch als Kund_innen oder Konsument_innen angesprochen. Mit dieser Entwicklung geht einher, dass unverzichtbare Aspekte wie Teilhabe an Daten, demokratische Kontrolle, Transparenz, Datenschutz und Qualitätssicherung von Daten nicht mehr gewährleistet sind.

In Europa sollte ein Alternativmodell der Datenökonomie aufgebaut werden, bei dem zahlreiche Akteure miteinander kooperieren und sowohl Daten bereitstellen als auch Daten nutzen können. Diese Form der Datenökonomie sollte mit Datensouveränität verbunden werden. Eine solche Möglichkeit bietet der Ansatz des Data Space (Datenraumarchitektur), in dem vielfältige Unternehmen und Einzelpersonen Daten bereitstellen, nutzen und miteinander austauschen. Die Daten werden kryptografisch gesichert und nur zertifizierte Anwender_innen können auf diese Daten für bestimmte Zwecke zugreifen.

KI-Systeme sollten immer im Anwendungskontext betrachtet und zertifiziert werden. Entscheidend ist dabei die Art der verwendeten Daten, die Qualitätssicherung der Daten und die Vorkehrungen zur Unterbindung von Fehlentwicklungen. Auf dieser Basis kann KI verantwortungsvoll, erklärbar, nachvollziehbar, transparent und gerecht eingesetzt werden.

Darüber hinaus sind in vielen Bereichen Investitionen erforderlich, um die gewünschte digitale Souveränität zu erreichen. So muss Europa in der technologischen Hardware unabhängiger werden. Im Bereich Computer- und Speicherleistungen sollten die begonnenen Investitionen weitergeführt werden, damit leistungsfähige Infrastrukturen auf- und ausgebaut werden.

7. Umgang mit Daten klären

Beim Einsatz von Daten in KI-Systemen sind grundlegende Entscheidungen zu treffen und es werden ethische und rechtliche Fragen von großer Bedeutung aufgeworfen. Die zentrale Frage lautet, wie mit Daten umgegangen werden soll und wer über sie verfügt. Diskutiert werden müssen z.B. folgende Fragen: Wem gehören die Daten, mit denen Maschinen lernen? Reichen die bestehenden gesetzlichen Regelungen aus? Wie können diskriminierende Auswirkungen von KI erkannt und verhindert werden? Inwieweit unterliegen Algorithmen einem Unternehmensgeheimnis oder sollten sie immer auch externen Kontrollinstanzen vorgelegt werden?

Bei der Diskussion über den Umgang mit Daten ist es erforderlich, in Bezug auf die Art der Daten und ihre Verwendung zu differenzieren. Ganz entscheidend ist dabei die Frage, ob es sich um personenbezogene oder nicht personenbezogene Daten handelt. Zu klären ist auch, wie weit die Transparenz und Offenlegungspflicht gehen sollte und ob Rohdaten oder aggregierte Daten betroffen sind. Daran schließen sich weitere Fragen an: Wem gegenüber soll die Offenlegungspflicht gelten? Soll z.B. ein zertifiziertes Unternehmen mit einem Schlüssel Zugriff darauf haben? Oder sollen bestimmte Daten grundsätzlich allgemein zugänglich sein (Open Data)? Auch muss entschieden werden, welche Datenarten für die jeweilige Nutzung und die Art des Zugangs in Frage kommen.

8. Schutz und Qualität der Daten sicherstellen

Die Konsequenzen von automatisiert getroffenen Entscheidungen eines KI-Systems können sehr weitreichend sein. Aufgrund von unvollständigen oder fehlerhaften Daten können Algorithmen falsche bzw. problematische Entscheidungen treffen. Deshalb muss die Qualität der in KI-Systemen verwendeten Daten gewährleistet werden: Die Daten müssen sicher, valide und zuverlässig sein und einen sorgfältigen Umgang mit personenbezogenen Daten einschließen. Zudem muss der Prozess der Qualitätssicherung nachvollziehbar dokumentiert werden. Der Staat sollte hier Standards setzen und für deren Einhaltung sorgen.

Notwendig sind klare gesetzliche Regelungen über die Zulässigkeit und die Grenzen der Datennutzung. Die Prinzipien des Datenschutzes in der europäischen Datenschutz-Grundverordnung (DSGVO) finden immer dann Anwendung, wenn KI mit persönlichen Daten arbeitet. Hier stellt sich die Frage, ob die gesetzlichen Regelungen der DSGVO ausreichen, um die Rechte der Bürger_innen zu schützen.

Um die Potenziale von KI für die Gesellschaft ausschöpfen zu können, muss die Forschung aber auch Zugang zu den erforderlichen Daten erhalten. Bisher ist es in vielen Bereichen, z.B. in der Medizin, noch eine große Herausforderung, die Nutzung von Daten im Kontext von maschinellem Lernen und KI zu organisieren. Hier gilt es, an angemessenen Lösungen weiterzuarbeiten. Ziel sollte es sein, das Individuum vor dem Missbrauch seiner Daten zu schützen und gleichzeitig die Möglichkeiten von KI dazu zu nutzen, die Daten für das Gemeinwohl und zur Bearbeitung gesellschaftlicher Herausforderungen einzusetzen.

9. Gesellschaftliche Debatte über Werte und Ziele führen

Wie bei jeder neuen technischen Entwicklung muss auch bei KI ein verantwortungsvoller Umgang mit den damit verbundenen Möglichkeiten erreicht werden. Deshalb bedarf es eines offenen gesellschaftlichen Diskurses über die Frage, wie die neuen technischen Möglichkeiten genutzt werden sollen und welche Regelungen und Grenzen notwendig sind, um Gefahren und Missbrauch entgegenzutreten. Dies kann nur durch eine breite öffentliche Debatte über die Werte und Ziele des Gemeinwesens erreicht werden, an dem sich Akteure aus Wissenschaft, Politik, Wirtschaft und Zivilgesellschaft beteiligen. Die Kernfrage lautet: „Welche Gesellschaft wollen wir und welchen Beitrag kann KI dazu leisten?“ Daran schließen sich weitere Grundsatzfragen an, die die gesamte Gesellschaft und ihre weitere Entwicklung betreffen, z.B.: Welche Formen der Mobilität und Kommunikation wollen wir in Zukunft? Welche Bedeutung sollen soziale Beziehungen in unserer Gesellschaft haben? Wie können die Fähigkeiten von Mensch und Maschine optimal zusammenwirken?

Die kontinuierliche Verständigung über gesellschaftliche Ziele und ethische Werte ist Kernbestandteil einer lebendigen Demokratie. Auf dieser Basis müssen politische Entscheidungen getroffen werden, um KI konstruktiv, transparent, gemeinwohlorientiert und menschenfreundlich zu gestalten.

EINFÜHRUNG

Der Prozess der Automatisierung war schon immer von großen Hoffnungen und Ängsten begleitet: Die einen sehen darin vor allem die Möglichkeit, wirtschaftliche Prosperität zu erwirken und der Menschheit Nutzen zu bringen, z.B. durch die Entlastung von stumpfsinniger und belastender Arbeit. Andere warnen vor den Gefahren eines ungehinderten technischen Fortschritts, etwa einer durch Automatisierung verursachten Massenarbeitslosigkeit oder einer drohenden Übermacht der Maschinen, und fordern eine politische Steuerung zum Wohl der Allgemeinheit.¹

Ein Teilgebiet der Automatisierung ist künstliche Intelligenz (KI), die durch die Digitalisierung erst möglich wurde. Hier wird intelligentes Problemlösungsverhalten automatisiert. In diesem Bereich wurden in den letzten Jahren enorme Entwicklungssprünge gemacht. KI ist zu einem großes Zukunftsthema geworden, da sie das Potenzial hat, Gesellschaft und Wirtschaft grundsätzlich zu verändern.

Im Zuge der digitalen Transformation ist davon auszugehen, dass menschliche und maschinelle Intelligenz zunehmend miteinander verknüpft werden und KI sich in vielen Bereichen durchsetzen wird, da sie zahlreiche Vorteile mit sich bringt. So erleichtert KI z.B. Prognosen, etwa durch die Voraussage des Ausfalls von Maschinen oder der wahrscheinlichen Entwicklung von Verkehrs- und Warenströmen. KI-Systeme können dabei helfen, den Verkehr fließend zu halten, Rohstoffkreisläufe zu optimieren, Krebszellen frühzeitig zu erkennen und den Produktionsprozess in einer Fabrik zu verbessern. Sie können Innovationsprozesse optimieren, bei der Zukunftsentwicklung von Städten assistieren und zur Gesundheitserhaltung vieler Menschen beitragen. Sie können Go spielen und Autos steuern. KI-Systeme haben aber keine Gefühle und Intelligenz im menschlichen Sinne, sie können weder Schmerz noch

¹ Martin Schwarz: „Zauberschlüssel zu einem Zukunftsparadies der Menschheit“. Automatisierungsdiskurse der 1950er- und 1960er-Jahre im deutsch-deutschen Vergleich (Dissertation an der TU Dresden). Dresden 2015, <https://nbn-resolving.org/urn:nbn:de:bsz:14-qucosa-191680> (30.03.2020).

Freude empfinden und auch nicht sterben, da sie Maschinen und keine Organismen sind. Künstliche Intelligenz ist nur dann wirklich „intelligent“, wenn sie durch humane Ziele und Werte gestaltet wird.²

Künstliche Intelligenz ist nur dann „intelligent“, wenn sie durch humane Ziele und Werte gestaltet wird.

KI ist der Anlass, (erneut) wichtige Fragen zu stellen, die die gesamte Gesellschaft und ihre weitere Entwicklung betreffen, z.B.: Welche Mobilitäts- und Kommunikationsformen wollen wir in Zukunft haben? Welche Bedeutung sol-

len soziale Beziehungen und persönliche Zuwendung haben, z.B. im Gesundheitssystem? Die damit verbundenen Ziele und Werte liegen jenseits maschineller Logik. Sie sind Hervorbringungen der menschlichen Kultur und Ausdrucksformen des Bewusstseins und der Empathie. KI fordert die Menschen heraus, sich ihrer gesellschaftlichen Ziele und Werte bewusst zu werden, damit die neuen technischen Möglichkeiten zum Nutzen der Menschheit konstruktiv und gemeinwohlorientiert eingesetzt werden können.

2

Zu diesem Abschnitt vgl. Zukunftsinstitut: 6 Thesen zur Künstlichen Intelligenz, <https://www.zukunftsinstitut.de/artikel/digitalisierung/6-thesen-zur-kuenstlichen-intelligenz/> (6.4.2020).

Was bedeutet ...?

Automatisierung bedeutet, dass die Funktionen des Produktionsprozesses, insbesondere Prozesssteuerungs- und -regelaufgaben vom Menschen auf künstliche Systeme übertragen werden. Automatisierung ist das Ergebnis des Automatisierens, d.h. des Einsatzes von Automaten, die als künstliche Systeme selbsttätig ein Programm befolgen und dabei aufgrund des Programms Entscheidungen zur Steuerung und Regelung von Prozessen treffen. Die Entscheidungen des Systems beruhen auf der Verknüpfung von Eingaben mit den jeweiligen Zuständen eines Systems und haben Aufgaben zur Folge. (1)

Künstliche Intelligenz (KI) oder Artificial Intelligence (AI) als fortgeschrittene Form der Automatisierung bezeichnet ein Teilgebiet der Informatik, in dem die Mechanismen des intelligenten menschlichen Verhaltens (Intelligenz) erforscht und nachgebaut werden. Dies geschieht im Wesentlichen durch Simulation mithilfe künstlicher Artefakte, meist mit Computerprogrammen. Eine Herausforderung ist dabei, dass „Intelligenz“ bisher noch nicht sehr gut definiert und verstanden ist. Im anwendungsbezogenen Sinn bedeutet KI, dass Maschinen Aufgaben ausführen, für die menschliche Intelligenzleistungen wie Lernen, Urteilen und Problemlösen nötig sind.

Maschinelles Lernen oder Machine Learning (ML) bedeutet, Computer so zu trainieren, dass sie anhand von Algorithmen in unstrukturierten Datensätzen Muster erkennen und aufgrund dieses „Wissens“ Entscheidungen treffen können. Ziel ist es, dass die Maschine aus Daten und Erfahrung lernt und ihre Aufgaben zunehmend besser ausführen kann.

Deep Learning ist ein Teilbereich des maschinellen Lernens und eine spezielle Methode der Informationsverarbeitung, bei der künstliche neuronale Netze genutzt und große Datenmengen analysiert werden. Der Computer greift dabei gleichzeitig auf Daten in mehreren Knotenebenen der neuronalen Netze zurück, um Zusammenhänge zu erkennen, Rückschlüsse zu ziehen und Vorhersagen sowie Entscheidungen zu treffen. Die Lernmethoden orientieren sich an der Funktionsweise des menschlichen Gehirns. Künstliche Intelligenz entsteht dann dadurch, dass das System das Erlernte immer wieder mit neuen

Inhalten verknüpft und (hinzu-)lernt. Dank selbstlernender Algorithmen kann die Maschine auch komplexe Probleme eigenständig lösen und ohne Anweisungen agieren. Deep Learning lehrt Maschinen zu lernen, d.h. sie werden in die Lage versetzt, selbstständig ihre Fähigkeiten zu verbessern, indem sie aus vorhandenen Daten und Informationen Muster extrahieren und klassifizieren. Die gewonnenen Erkenntnisse lassen sich mit Daten korrelieren und in einem weiteren Kontext verknüpfen. Schließlich kann die Maschine Entscheidungen auf Basis der Verknüpfungen treffen. (2)

Bei der Anwendung wird Deep Learning im Rahmen künstlicher Intelligenz häufig für die Gesichts-, Objekt- oder Spracherkennung eingesetzt. Bei der Spracherkennung können die Systeme z.B. ihren Wortschatz selbstständig mit neuen Wörtern erweitern (etwa der intelligente Sprachassistent Siri von Apple), gesprochene Texte übersetzen, Computerspiele „intelligenter“ machen, das autonome Fahren ermöglichen oder Kundenverhalten auf der Basis von Daten eines CRM (Customer-Relationship-Management)-Systems vorhersagen. (3)

Quellen: (1) Prof. Dr. Kai-Ingo Voigt: Automatisierung, 19.2.2018, Gabler Wirtschaftslexikon, <https://wirtschaftslexikon.gabler.de/definition/automatisierung-27138/version-250801>; (2) Spektrum.de: Künstliche Intelligenz. Lexikon der Neurowissenschaft, <https://www.spektrum.de/lexikon/neurowissenschaft/kuenstliche-intelligenz/6810>; (3) Stefan Luber: Was ist Deep Learning? In: Definitionen, Big Data Insider, <https://www.bigdata-insider.de/was-ist-deep-learning-a-603129/> (15.08.2019).

Aufgrund der großen Bedeutung der Künstlichen Intelligenz für Gesellschaft und Wirtschaft fördert die Bundesregierung gezielt die Forschung und Anwendung von KI mit einem Programm, um Deutschland in diesem Bereich erheblich zu stärken. Die Strategie Künstliche Intelligenz der Bundesregierung (KI-Strategie) wurde nach Durchführung eines deutschlandweiten Online-Konsultationsverfahrens unter gemeinsamer Federführung des Bundesministeriums für Bildung und Forschung (BMBF), des Bundesministeriums für Wirtschaft und Energie (BMWi) und des Bundesministeriums für Arbeit und Soziales (BMAS) erstellt. Ziel ist es, einen Rahmen für die ganzheitliche politische Gestaltung der weiteren Entwicklung und Anwendung Künstlicher Intelligenz in Deutschland zu setzen.

Strategie Künstliche Intelligenz der Bundesregierung (2018)

Mit der KI-Strategie verfolgt die Bundesregierung drei wesentliche Ziele:

1. Deutschland soll zu einem **weltweit führenden Standort in Forschung und Anwendung von KI** werden („KI Made in Germany“), um zur Sicherung der internationalen Wettbewerbsfähigkeit beizutragen. Folgende Maßnahmen sollen umgesetzt werden:

- Überregionaler Ausbau und Schaffung neuer Kompetenzzentren für KI-Forschung zu einem nationalen Netzwerk von mindestens zwölf Zentren und Anwendungshubs,
- Auflegen eines Programms zur wissenschaftlichen Nachwuchsförderung und Lehre im Bereich KI und Schaffung von mindestens hundert zusätzlichen KI-Professuren,
- Aufbau eines deutsch-französischen Forschungs- und Innovationsnetzwerkes („virtuelles Zentrum“),
- KI als Schwerpunkt der Agentur für Sprunginnovationen,
- Bildung eines europäischen Innovationsclusters zu KI,
- Ausweitung der KI-spezifischen Unterstützung von mittelständischen Unternehmen durch „KI-Trainer_innen“,
- Unterstützung von Unternehmen bei der Einrichtung von Testfeldern,
- Verdopplung der Haushaltsmittel in 2019 für EXIST, dem Programm für Existenzgründungen aus der Wissenschaft,
- Ausbau von öffentlichen Förderangeboten im Bereich Wagniskapital und Venture Debt und Start einer Tech Growth Fund Initiative,
- Ausbau von Angeboten zur ganzheitlichen Beratung und Förderung von Gründungen,
- Verbesserung von Anreizen und Rahmenbedingungen für das freiwillige, datenschutzkonforme Teilen von Daten sowie Aufbau einer vertrauenswürdigen Daten- und Analyseinfrastruktur (einschließlich Cloud-Plattform mit skalierbarer Speicher- und Rechenkapazität).

2. KI soll verantwortungsvoll und gemeinwohlorientiert entwickelt und genutzt werden:

- Einrichtung eines deutschen Observatoriums für Künstliche Intelligenz und Engagement für den Aufbau entsprechender Observatorien auf europäischer und internationaler Ebene,
- Organisation eines europäischen und transatlantischen Dialogs zum menschenzentrierten Einsatz von KI in der Arbeitswelt,
- Entwicklung eines breitenwirksamen Instrumentariums zur Förderung der Kompetenzen von Erwerbstätigen im Rahmen einer Nationalen Weiterbildungsstrategie,
- Weiterentwicklung der Fachkräftestrategie im Hinblick auf den digitalen Wandel und die neuen KI-Technologien auf der analytischen Grundlage eines neuen Fachkräftemonitorings,
- Sicherung der betrieblichen Mitbestimmungsmöglichkeiten bei der Einführung und Anwendung von KI,
- Förderung von betrieblichen Experimentierräumen zu KI-Anwendungen in der Arbeitswelt,
- Förderung von KI-Anwendungen zum Nutzen von Umwelt und Klima und Entwicklung von dafür geeigneten Bewertungsgrundlagen, Anstoß von fünfzig Leuchtturmanwendungen.

3. Im Rahmen eines breiten gesellschaftlichen Dialogs und einer aktiven politischen Gestaltung soll KI ethisch, rechtlich, kulturell und institutionell derart in die Gesellschaft eingebettet werden, dass gesellschaftliche Grundwerte und individuelle Grundrechte gewahrt bleiben und die Technologie der Gesellschaft und dem Menschen dient:

- Einrichten eines Runden Tisches mit Datenschutzaufsichtsbehörden und Wirtschaftsverbänden, um gemeinsam Leitlinien für eine datenschutzrechtskonforme Entwicklung und Anwendung von KI-Systemen zu erarbeiten sowie Aufbereitung von Best-Practice-Anwendungsbeispielen,
- Förderung der Entwicklung von innovativen Anwendungen, die die Selbstbestimmung, die soziale und kulturelle Teilhabe sowie den Schutz der Privatsphäre der Bürgerinnen und Bürger unterstützen,
- Förderung von Aufklärung und multidisziplinärer sozialer Technikgestaltung in der Breite durch einen „ZukunftsFonds Digitale Arbeit und Gesellschaft“,

- Weiterentwicklung der Plattform Lernende Systeme zu einer Plattform für Künstliche Intelligenz, in der ein Austausch zwischen Politik, Wissenschaft und Wirtschaft mit der Zivilgesellschaft organisiert wird.

Die vorliegende KI-Strategie versteht sich als Handlungsrahmen der Bundesregierung, die Anfang 2020 je nach Diskussionsstand und Erfordernissen weiterentwickelt und beim Fortschreiben den neuesten Entwicklungen und Bedarfen angepasst werden soll.

Für die Umsetzung der Strategie will der Bund zwischen 2019 und 2025 insgesamt etwa drei Mrd. Euro zur Verfügung stellen. Die Hebelwirkung dieses Engagements auf Wirtschaft, Wissenschaft und Länder soll mindestens zur Verdoppelung dieser Mittel führen. Mit dem Bundeshaushalt 2019 hat der Bund in einem ersten Schritt insgesamt 500 Mio. zur Verstärkung der KI-Strategie für 2019 und die Folgejahre bis 2023 zur Verfügung gestellt:

- Maßnahmen für den KI-Transfer in die Anwendung in der Wirtschaft (ca. 230 Mio. Euro),
- Förderung der Forschung allgemein und des wissenschaftlichen Nachwuchses (ca. 190 Mio. Euro),
- Projekte in den Bereichen gesellschaftlicher Dialog und Partizipation, Technikfolgenabschätzung und Ordnungsrahmen sowie Förderung betrieblicher Qualifikationsmaßnahmen (ca. 55 Mio. Euro).

Quellen: Strategie Künstliche Intelligenz der Bundesregierung. Stand: November 2018, https://www.bmbf.de/files/Nationale_KI-Strategie.pdf; Stefan Krempel: heise online, 24.5.2019, <https://www.heise.de/newsticker/meldung/KI-Strategie-Bundesregierung-verteilt-die-erste-halbe-Milliarde-4431505.html> (6.4.2020).

Politische Initiativen im Bereich KI. Aufgrund der großen Bedeutung des Themas KI für die gesellschaftliche und wirtschaftliche Entwicklung beschäftigen sich derzeit verschiedene Gremien in Deutschland mit diesen Fragen:

- Der Bundestag hat eine **Enquete-Kommission** „Künstliche Intelligenz – Gesellschaftliche Verantwortung und wirtschaftliche, soziale und ökologische Potenziale“ eingesetzt, die sich am 27. September 2018 konstituiert hat. Dem Gremium gehören 19 Bundestagsab-

geordnete und 19 externe Sachverständige an. Die Kommission soll den zukünftigen Einfluss der Künstlichen Intelligenz (KI) auf das (Zusammen-)Leben, die deutsche Wirtschaft und die zukünftige Arbeitswelt untersuchen. Erörtert werden sowohl die Chancen als auch die Herausforderungen der KI für Gesellschaft, Staat und Wirtschaft. Zur Diskussion stehen eine Vielzahl technischer, rechtlicher und ethischer Fragen. Zum Auftrag der Enquete-Kommission gehört, auf Basis ihrer Untersuchungsergebnisse den staatlichen Handlungsbedarf auf nationaler, europäischer und internationaler Ebene zu identifizieren und zu beschreiben, um die Chancen der KI wirtschaftlich und gesellschaftlich nutzbar zu machen und ihre Risiken zu minimieren.³

- Der 2018 neu geschaffene **Digitalrat** soll die Bundesregierung bei der weiteren Gestaltung ihrer Netzpolitik beraten. Im Digitalrat arbeiten unabhängige Expert_innen aus Wissenschaft, Forschung und Wirtschaft zusammen, die verschiedene Erfahrungshintergründe haben und das große Spektrum der Digitalszene abbilden sollen. Die Mitglieder des Rats sollen theoretische Erkenntnisse, Praxiserfahrung und Innovationen verbinden und die Bundesregierung mit ihrer Fachexpertise unterstützen – dabei soll der Rat auch unbequem sein und die Politik „antreiben“.⁴
- Die KI-Strategie ist Teil der **Digitalisierungsstrategie**⁵ der Bundesregierung, die dazu beitragen soll, den digitalen Wandel zu gestalten und Deutschland auf die Zukunft im digitalen Zeitalter bestmöglich vorzubereiten. In der Umsetzungsstrategie zur Gestaltung des digitalen Wandels wurden fünf entscheidende Handlungsfelder festgelegt: Digitale Kompetenz, Infrastruktur und Ausstattung, Innovation und digitale Transformation, Gesellschaft im digitalen Wandel sowie Moderner Staat.
- Die aus 16 Expert_innen bestehende **Datenethikkommission** beschäftigt sich mit ethischen und rechtlichen Fragen beim Einsatz

3 Deutscher Bundestag: Enquete-Kommission „Künstliche Intelligenz“ hat sich konstituiert, <https://www.bundestag.de/dokumente/textarchiv/2018/kw39-pa-enquete-kuenstliche-intelligenz-567956>, https://www.bundestag.de/ausschuesse/weitere_gremien/enquete_ki (6.4.2020).

4 Die Bundesregierung: der Digitalrat – Experten, die uns antreiben, <https://www.bundesregierung.de/breg-de/themen/digitalisierung/der-digitalrat-experten-die-uns-antreiben-1504866> (6.4.2020).

5 Digitalisierungsstrategie der Bundesregierung, <https://www.bildung-forschung.digital/de/digitalisierungsstrategie-der-bundesregierung-2529.html> (6.4.2020).

von Algorithmen und KI. Das Gremium unter Federführung des Bundesministeriums des Innern, für Bau und Heimat und des Bundesministeriums der Justiz und für Verbraucherschutz sollte die Bundesregierung auf der Basis von wissenschaftlicher und technischer Expertise beraten und „ethische Leitlinien für den Schutz des Einzelnen, die Wahrung des gesellschaftlichen Zusammenlebens und die Sicherung und Förderung des Wohlstands im Informationszeitalter entwickeln“.⁶ Die Kommission tagte zum ersten Mal am 5. September 2018 und übergab im Oktober 2019 in ihrem Abschlussgutachten ihre Handlungsempfehlungen zum Umgang mit Daten und algorithmischen Systemen an die Bundesregierung.⁷

- Am 11. Oktober 2018 nahm eine neue, interdisziplinäre und agil arbeitende Abteilung im Bundesministerium für Arbeit und Soziales (BMAS) ihre Tätigkeit auf: Die **Denkfabrik Digitale Arbeitsgesellschaft** verbindet Funktionen und Arbeitsweisen eines klassischen Think Tanks mit denen eines zeitgenössischen Future Labs. Sie hat die Aufgabe, Szenarien zur Zukunft der Arbeit im digitalen Zeitalter zu entwickeln und neue Handlungsfelder für die Politik zu identifizieren. Sie beteiligt sich an der Digitalisierungsstrategie der Bundesregierung und untersucht in einem neuen, einjährigen Programmschwerpunkt die Verschiebung von Macht- und Kooperationsverhältnissen in der digitalen Transformation.⁸ Am 3. März 2020 wurde das deutsche „**KI-Observatorium**“ eröffnet, das als eine Art TÜV für KI im Wirtschafts- und Arbeitsleben fungieren soll. Es ist zunächst als Einheit im BMAS angesiedelt und soll später zu einem Bundesinstitut werden. Aufgabe des Observatoriums ist es, Chancen und Risiken der künstlichen Intelligenz zu bewerten und die Politik dabei zu unterstützen, KI politisch zu steuern und zu gestalten. Ziel ist ein europaweites Netz kooperierender KI-Bewertungsstellen. Das deutsche Observatorium hat jedoch eine besondere

6 Bundesministerium des Innern, für Bau und Heimat, <https://www.bmi.bund.de/DE/themen/it-und-digitalpolitik/datenethikkommission/datenethikkommission-node.html> (6.4.2020).

7 Zu den zentralen Handlungsempfehlungen gehören u.a. ein risikoadaptiertes Regulierungssystem für den Einsatz von algorithmischen Systemen (mit nach Schädigungspotenzial abgestufter Regulierung), die Schaffung eines bundesweiten „Kompetenzzentrums Algorithmische Systeme“, die Etablierung einer EU-Verordnung mit Grundanforderungen an die Zulässigkeit von algorithmischen Systemen sowie spezifische rechtliche Vorgaben für persönlichkeitsensible Profilbildungen von Qualitätsanforderungen bis hin zu absoluten Grenzen. Vgl. Pressemitteilung des Bundesministeriums des Innern, für Bau und Heimat, 23.10.2019, <https://www.bmi.bund.de/SharedDocs/pressemitteilungen/DE/2019/10/datenethikkommission-okt-2019.html> (6.4.2020).

8 Vgl. <https://www.denkfabrik-bmas.de/> (6.4.2020).

Ausrichtung: Die betriebliche Mitbestimmung wird im BMAS als Erfolgsfaktor bei der Einführung von künstlicher Intelligenz in Produktionsprozesse betrachtet, da sie das Vertrauen der Beschäftigten stärke und dabei helfe, die neue Technologie zu etablieren. Deshalb will das BMAS 2020 auch eine Novelle des Betriebsverfassungsgesetzes auf den Weg bringen, um die Mitbestimmung auf die künstliche Intelligenz auszudehnen.⁹

- Die **Europäische Gruppe für Ethik der Naturwissenschaften und der Neuen Technologien** hat im März 2018 eine Erklärung veröffentlicht, in der die Einleitung eines Prozesses gefordert wird, „den Weg für die Entwicklung eines gemeinsamen, international anerkannten ethischen und rechtlichen Rahmens für die Konstruktion, Produktion, Verwendung und Steuerung von künstlicher Intelligenz, Robotik und ‚autonomen‘ Systemen bereiten“ soll.¹⁰ Seitdem werden Leitlinien zur Ethik in der künstlichen Intelligenz ausgearbeitet.

Themen der Debatte. Dieser Publikation liegen Beiträge einer Konferenz der Friedrich-Ebert-Stiftung zugrunde, in der der Schwerpunkt auf KI in Wissenschaft und Forschung gelegt wurde.¹¹ Thematisiert wurden zum einen die Chancen und Potenziale von KI-Systemen, wie z.B. die Möglichkeit, Kerninformationen aus nicht strukturierten Texten herauszukristallisieren, was große Fortschritte in der Forschung mit sich bringen könnte. Zum anderen wurden die Gefahren und Risiken angesprochen, z.B.:

- Wie soll mit den Datenmonopolen einiger Digitalfirmen (Facebook, Google, Apple ...) umgegangen werden, die mit ihren globalen Digitalplattformen ganze Wirtschaftsbereiche dominieren?
- Welche Risiken sind mit der Datenerhebung verbunden, die für die Entwicklung von KI notwendig ist, und wie kann ihnen wirksam begegnet werden?

⁹ Vgl. Hendrik Munsberg: TÜV für künstliche Intelligenz kommt. In: Süddeutsche Zeitung online, 12.11.2019, <https://www.sueddeutsche.de/wirtschaft/ki-observatorium-tuev-arbeitsministerium-1.4676937> (6.4.2020).

¹⁰ Europäische Kommission, Generaldirektion Forschung und Innovation, Europäische Gruppe für Ethik der Naturwissenschaften und der Neuen Technologien: Erklärung zu künstlicher Intelligenz, Robotik und „autonomen“ Systemen, Brüssel, 9. März 2018, https://ec.europa.eu/research/eye/pdf/eye_ai_statement_2018_de.pdf (6.4.2020).

¹¹ Konferenz „Bit für Bit in die Zukunft? Künstliche Intelligenz in Wissenschaft und Forschung, 7. März 2019 in Berlin.

- Wie kann zukünftig Datensicherheit gewährleistet werden?
- Welche rechtlichen und ethischen Aspekte müssen bei der Entwicklung von KI-Systemen berücksichtigt werden?
- Welche gesellschaftlichen Erwartungen sind mit KI verbunden?

Die Entwicklung von KI hat darüber hinaus auch einen internationalen Kontext: Dabei verfolgen die Länder unterschiedliche Strategien und setzen andere Schwerpunkte – von militärisch geprägter KI-Forschung in den USA bis hin zu einem starken Fokus auf Unterhaltungstechnologie in Japan oder Südkorea. Hier stellt sich die Frage, welchen Weg Deutschland gehen soll und wodurch sich „KI Made in Germany“ auszeichnen sollte. Wäre es z.B. ein guter Weg, sich auf medizinische Forschung oder nachhaltige Technologien zu konzentrieren? Welche Rolle spielt dabei Europa und welche Kooperationsmöglichkeiten hat Deutschland mit anderen europäischen Ländern, z.B. mit Frankreich?

Gegenwärtig wird die Entwicklung der KI-Technologien von vielen noch als bedrohlich wahrgenommen. Unklar ist z. B., wie eine Intelligente Technologie kontrolliert werden kann und wer die Verantwortung trägt, wenn autonome Systeme Fehler machen. Ängste entstehen aber auch dadurch, dass Menschen durch Roboter ersetzt werden könnten und dadurch ihren Arbeitsplatz verlieren. Deshalb standen bei der Konferenz auch Fragen der politischen Regulierung und des Datenschutzes, speziell in Bezug auf Forschung und Wissenschaft, im Austausch von Expert_innen und Politik im Vordergrund.

Wodurch sollte sich „KI Made in Germany“ auszeichnen?

KÜNSTLICHE INTELLIGENZ: GROSSE CHANCEN, GROSSE RISIKEN

Einen historischen Abriss über die Geschichte der Künstlichen Intelligenz gab Prof. Dr. Ute Schmid, die an der Fakultät Wirtschaftsinformatik und Angewandte Informatik an der Otto-Friedrich Universität Bamberg forscht und lehrt.

Anfänge der KI. Schmid machte deutlich, dass KI in den letzten Jahren zunehmend als Buzz Word in Debatten erscheint, aber keineswegs eine neue Erfindung ist, sondern schon eine längere Geschichte hat. Der Ausgangspunkt der Künstlichen Intelligenz liegt im Jahr 1956 bei der Dartmouth Conference im US-Bundesstaat New Hampshire in den USA. Dabei handelte es sich um ein Forschungsprojekt zur KI, das der Informatiker John McCarthy¹², der KI-Forscher Marvin Minsky, IBM-Mitarbeiter Nathaniel Rochester und der Informationstheoretiker Claude Shannon planten und mit weiteren Wissenschaftlern durchführten. Ihre Grundannahme war, dass jeder Aspekt von Intelligenz so präzise beschrieben werden kann, dass ein Computer ihn simulieren kann. Die Konferenz gilt als Geburtsstunde der KI als akademisches Fachgebiet und als Ursprung des Namens „Artificial Intelligence“, der später ins Deutsche („künstliche Intelligenz“) übersetzt wurde.

Phasen der KI. Schmid erläuterte die verschiedenen Phasen der KI-Forschung, die immer wieder von überzogenen Versprechungen und Erwartungen und anschließender Ernüchterung und Enttäuschung geprägt waren. In der ersten Phase ab 1956 befasste sich die KI-Forschung mit Aspekten von intelligentem Verhalten, vor allem in den Bereichen Spiele, maschinelles Lernen, Problemlösung, Planen und Sprachverstehen. Bereits in dieser frühen Zeit seien sehr gute Algorithmen entwickelt worden, die bis heute ihren Platz in Anwendungen haben,

¹² McCarthy stellte kurze Zeit später die eigene Programmiersprache Lisp für die Verarbeitung symbolischer Strukturen vor, die in den folgenden Jahrzehnten vor allem in den USA die Standardsprache für KI-Anwendungen werden sollte. Vgl. Klaus Manhart: Eine kleine Geschichte der Künstlichen Intelligenz, 17.1.2018, <https://www.computerwoche.de/a/eine-kleine-geschichte-der-kuenstlichen-intelligenz,33305372> (6.4.2020).

meinte Schmid. Auch damals seien die Erwartungen schon überzogen gewesen: So glaubte z.B. Herbert A. Simon¹³ im Jahr 1957, dass in den nächsten zehn Jahren ein Computer Schachweltmeister werden könnte. Tatsächlich gelang dies aber erst vierzig Jahre später, als 1997 der von IBM entwickelte Schachcomputer Deep Blue den damals amtierenden Schachweltmeister Garri Kasparow in einem Wettkampf besiegte. Zahlreiche weitere Fehlschläge (z.B. bei der maschinellen Übersetzung 1966) und ein insgesamt langsamer Fortschritt führten dann zum ersten „KI-Winter“ (1974 bis 1980). Als nächster Hype folgten Expertensysteme: Da bereits sehr gute Inferenzformalismen¹⁴ und KI-Sprachen (Lisp, Prolog) entwickelt worden waren, sei man der Ansicht gewesen, dass die Maschinen jetzt nur noch mit Expertenwissen gefüttert werden müssten, erläuterte Schmid. Diese Annahme habe sich jedoch als naiv erwiesen, da Wissen nicht so leicht und schnell in eine Maschine integriert werden kann, wenn es explizit modelliert wird – im Unterschied zum allmählichen maschinellen Lernen. Dann folgte der nächste KI-Winter (1987 bis 1993). Erst nach einer weiteren Forschungs- und Entwicklungsphase und einem erneuten KI-Winter (2000 bis 2008) begannen 2008 die großen Erfolge von Machine Learning. Fortschritte gab es sowohl im Bereich Bildverarbeitung bzw. der Objekterkennung auf Bildern (2012 erkannte Google Brain das Bild einer Katze) als auch im Bereich Spiele. Zudem folgte der Big Bang beim Deep Learning¹⁵. Aktuell befinde man sich im nächsten Hype und es sei nur eine Frage der Zeit, wann der nächste KI-Winter komme, meinte Schmid, die vor überzogenen Erwartungen warnte: „Die KI kann sicher viel leisten, aber man sollte rational bleiben und nicht dem Hype folgen.“

13 Der US-amerikanische Sozialwissenschaftler, Psychologe und Wirtschaftsnobelpreisträger Herbert A. Simon (1916-2001) hatte 1956 gemeinsam mit dem Informatiker und Kognitionspsychologen Allen Newell (1927-1992) das Programm Logical Theorist entwickelt, das erstmals logische Theoreme beweisen konnte. Dies gilt als Meilenstein der Künstlichen Intelligenz, da deutlich wurde, dass Programme zu Aktionen fähig sind, für die ein Mensch Intelligenz braucht. Ein weiterer wichtiger Schritt hin zum maschinellen Problemlösen war die ebenfalls mit Newell entwickelte Software General Problem Solver (allgemeiner Problemlöser). Vgl. Digitale Pioniere: Herbert A. Simon. Problemlöser. In: Tagesspiegel, 5.4.2026, <https://www.tagesspiegel.de/wissen/digitale-pioniere-39-herbert-a-simon-problemloeser/13400994.html> (6.4.2020).

14 Inferenz ist eine Handlung oder ein Prozess, der logische Schlussfolgerungen aus bekannten oder angenommenen Prämissen ableitet. In der KI wird darunter die Verwendung eines trainierten Modells verstanden, um Vorhersagen über Daten zu machen. Vgl. Sebastian Gerstl: KI-Definition. 6 Fachbegriffe der Künstlichen Intelligenz erklärt, in: Elektronik Praxis, 22.5.2018, <https://www.elektronikpraxis.vogel.de/ki-definition-6-fachbegriffe-der-kuenstlichen-intelligenz-erklart-a-717094/index2.html> (6.4.2020).

15 Siehe Begriffserklärung „Deep Learning“, Info-Kasten in der Einführung dieser Publikation.

Gegenstand der KI. Schmid machte darauf aufmerksam, dass es zahlreiche Definitionen von KI gibt. Zum Teil werde explizit vermieden, den Begriff „Intelligenz“ in der Definition zu verwenden. Aus Sicht von Marvin Minsky, einem der Gründungsväter der KI, befasst sich KI-Forschung mit Informatik-Problemen, die bislang ungelöst sind. Nach Auffassung der Informatikerin Elaine Rich erforscht KI-Wissenschaft, wie Computer dazu befähigt werden können, Dinge zu tun, die Menschen im Moment noch besser können. Konsens besteht nach Schmid jedoch darüber, dass KI ein Teilgebiet der Informatik ist, das viele interdisziplinäre Beziehungen aufweist, vor allem zur Psychologie (Kognitive Systeme), zur Philosophie des Geistes und zur Sprachwissenschaft. Es gelte: „KI ist das, was KI-Forscherinnen und -Forscher machen.“

„Die KI kann sicher viel leisten, aber man sollte rational bleiben.“

Starke und schwache KI. Zu unterscheiden sind zwei Formen von KI: die „schwache“ und die „starke“ KI. Die schwache KI fokussiert auf die Entwicklung von informatischen Lösungen für spezielle Bereiche mit komplexen Anforderungen, etwa Suchverfahren mit verschiedenen Heuristiken¹⁶. In diesem Bereich sind die meisten KI-Forscher_innen tätig. In der Öffentlichkeit wird aber vorrangig die „starke KI“ diskutiert, die darauf abzielt, natürliche Intelligenz zu verstehen und nachzubauen. Anspruch der starken KI ist die Entwicklung von Systemen, die nicht nur reagieren, sondern auch aus eigenem Antrieb intelligent und flexibel handeln können. Grundannahme ist, dass Computer/Roboter auf die gleiche Art intelligent sein könnten wie Menschen.¹⁷

„Maschinen werden Fehler weniger verzeihen als Menschen.“

¹⁶ Heuristik bezeichnet eine Vorgehensweise zur Lösung von mathematischen Problemen, die auf der Basis von Erfahrungen oder Urteilsvermögen zu einer guten Lösung führt, die aber nicht unbedingt optimal ist. Vgl. Jean-Paul Thommen: Heuristik. In: Gabler Wirtschaftslexikon, <https://wirtschaftslexikon.gabler.de/definition/heuristik-34474> (6.4.2020).

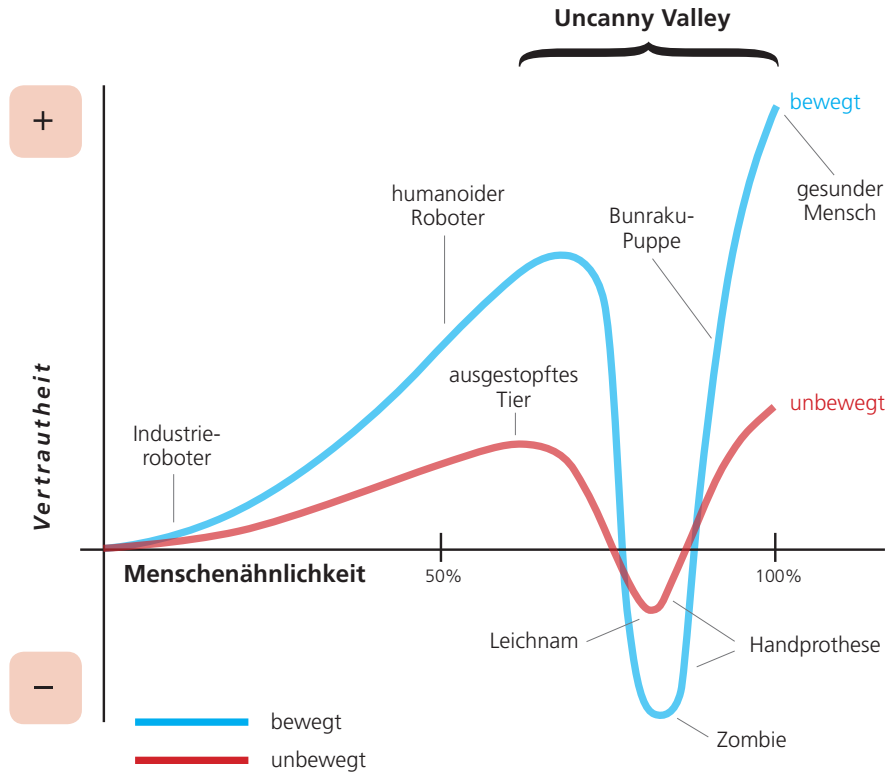
¹⁷ Im Mittelpunkt der schwachen KI stehen Systeme, die sich auf die Lösung konkreter Anwendungsprobleme konzentrieren, z.B. Experten- und Navigationssysteme, Sprach- und Zeichenerkennung oder Korrekturvorschläge bei Suchprozessen. Dagegen wird mit der starken KI das Ziel verfolgt, eine allgemeine Intelligenz zu schaffen, die der des Menschen gleicht oder diese sogar übertrifft. Eine wirkliche Intelligenz müsste folgende Fähigkeiten beherrschen: Logisches Denken, Treffen von Entscheidungen bei Unsicherheit, Planen, Lernen, Kommunikation in natürlicher Sprache sowie Einsatz dieser Fähigkeiten zum Erreichen eines Ziels. Darüber hinaus wird starke KI noch mit Bewusstsein, Selbsterkenntnis/Eigenwahrnehmung, Empfindungsvermögen und Weisheit assoziiert. Vgl. Schwache KI und Starke KI, in: Künstliche Intelligenz. Informatik und Gesellschaft 2008/2009, http://www.informatik.uni-oldenburg.de/~iug08/ki/Grundlagen_Starke_KI_vs._Schwache_KI.html (6.4.2020).

Autonomes Fahren. In öffentlichen Debatten steht auch immer wieder die Frage des autonomen Fahrens¹⁸ im Mittelpunkt. Kritisch wird z.B. gesehen, dass ein selbstfahrendes Auto in eine Verkehrssituation geraten könnte, die eine ethisch schwierige Entscheidung erforderlich macht, wenn Menschenleben gefährdet werden. So könnte z. B. die Frage entstehen, welches Menschenleben mehr „wert ist“: das Leben eines alten Menschen oder einer schwangeren Frau. Schmid meinte, es sei zwar grundsätzlich möglich, eine Geschlechts- und Alterserkennung in solche Systeme einzubauen, sodass z.B. ein autonomes Fahrzeug Fußgänger_innen erkennen kann. Doch halte sie es für sehr unwahrscheinlich, dass die Entwicklung in diese Richtung gehe, weil gar nicht klar sei, wie dann mit dieser Information umgegangen werden sollte. Aus Schmid's Sicht ist es für praktische Anwendungen eher nicht erstrebenswert, KI-Systeme entwickeln zu wollen, die einen eigenen Willen haben und bewusst Entscheidungen treffen können.

Fehler von KI-Systemen. Immer wieder werde auf die Fehleranfälligkeit von KI-Systemen verwiesen, was die Menschen verunsichere: „Maschinen werden Fehler weniger verzeihen als Menschen“, sagte Schmid. Dies liege auch daran, dass Maschinen oft Fehler machen, die wenig menschenähnlich sind. Menschen würden auf eine bestimmte Art und Weise Fehler machen, und in der Interaktion könnten diese dann meist erklären, wie der jeweilige Fehler passieren konnte. Wenn ein Computer oder ein Roboter einen Fehler macht, sei das anders. Dann offenbare sich, dass eine solche Maschine letztlich ein „Fachidiot“ ist. Interessant sei in diesem Zusammenhang auch das Phänomen des „Uncanny

18 Der Begriff „autonomes Fahren“ bezeichnet das „selbstständige, zielgerichtete Fahren eines Fahrzeugs im realen Verkehr, ohne Eingriff des Fahrers“. Vgl. Daimler, <https://www.daimler.com/innovation/case/autonomous/rechtlicher-rahmen.html> (6.4.2020). Dagegen bedeutet „automatisiertes Fahren“, dass Fahrer_innen durch Assistenz- und teilautomatisierte Systeme (wie z.B. Stop&Go Pilot und Aktiver Spurwechsel-Assistent) unterstützt werden, den Menschen aber nicht ersetzen. Der Unterschied zwischen automatisiertem und autonomem Fahren ist auch juristisch von Bedeutung. In Deutschland sind am 21. Juni 2017 neue Regeln zum automatisierten Fahren in Kraft getreten. Das Gesetz schafft die rechtlichen Voraussetzungen für hoch- und voll automatisierte Systeme, die die Fahrzeugsteuerung komplett übernehmen, auch wenn die (mit) fahrende Person übernahmebereit bleiben muss. Allerdings regelt das neue Gesetz nicht das autonome Fahren, bei dem es nur noch Passagiere gibt. Im Jahr 2016 hat die Bundesregierung eine Ethikkommission aus Expert_innen eingesetzt, die sich mit rechtlichen und ethischen Fragen beim autonomen Fahren auseinandersetzt. Im Juni 2017 verabschiedete die Kommission einen Abschlussbericht mit insgesamt zwanzig ethischen Regeln, darunter die Festlegung, dass der Schutz des Menschen immer Vorrang haben muss. Vgl. ebd., Bundesministerium für Verkehr und digitale Infrastruktur: Bericht der Ethik-Kommission Automatisiertes und vernetztes Fahren. Juni 2017, https://www.bmvi.de/SharedDocs/DE/Publikationen/DG/bericht-der-ethik-kommission.pdf?__blob=publicationFile (6.4.2020).

Valley“: In Studien sei festgestellt worden, dass eine Akzeptanzlücke bei Robotern besteht, die den Menschen sehr ähneln. Sie machten den Menschen Angst und lösten starke Irritationen aus.¹⁹



Quelle: https://de.wikipedia.org/wiki/Uncanny_Valley#/media/Datei:Mori_Uncanny_Valley_de.svg

¹⁹ In der Robotik bezeichnet der Begriff Uncanny Valley („unheimliches Tal“) oder Akzeptanzlücke den messbaren Effekt, dass die menschliche Akzeptanz für Roboter oder Avatare etc. schlagartig abfällt, wenn diese dem Menschen zu sehr ähneln. Das besagte „Tal“ bezeichnet einen starken Einbruch in einem Kurvendiagramm, das den Verlauf der Ablehnung bzw. Akzeptanz von Robotern darstellt. Vgl. Richard Watson: Uncanny Valley, SpringerLink, https://link.springer.com/chapter/10.1007/978-3-642-40744-4_35 (6.4.2020).

Begriff der Intelligenz. Schmid erläuterte, dass der Intelligenzbegriff der Psychologie auf der kognitiven Leistungsfähigkeit des Menschen basiert, wie sie auch in Intelligenztests geprüft wird und schlussfolgerndes Denken (induktiv, deduktiv, abduktiv, mit Analogien) voraussetzt.²⁰ Landläufig werde Intelligenz aus Verhaltensweisen abgeleitet, bei denen das Vorhandensein von Intelligenz vorausgesetzt wird, zum Beispiel wenn ein Mensch Schach spielt oder Mathematik studiert. So werde z.B. angenommen, dass ein Mensch, der Schach spielen kann, in der Regel auch eine Katze auf einem Bild erkennen kann. Dadurch übertrage der Mensch seine Erwartungen an KI-Systeme, was aber letztlich nicht zutreffe. Gegenwärtig könnten Menschen noch einiges besser als Maschinen, etwa komplizierte Konstruktionen aus Bauklötzen errichten, Witze verstehen, Fahrrad fahren oder empathisch sein. Ein grundsätzliches Problem besteht nach Schmid darin, dass Intelligenz immer nur eine Zuschreibung aufgrund von Verhalten ist, wie z.B. der Turing-Test zeige. Intelligenz werde hier aber nur vermutet, man wisse es nicht. Unverzichtbar sei etwas anderes: „Intelligenz basiert auf Intentionalität“, sagte Schmid.

*„Intelligenz basiert auf
Intentionalität.“*

20

In der Erkenntnistheorie werden verschiedene Arten der Schlussfolgerung unterschieden: deduktiv (Ableitung von einer allgemeinen Aussage/Regel auf einen Einzelfall), induktiv (Folgerung vom Einzelfall auf Allgemeines und Gesetzmäßigkeiten), abduktiv (hypothesen- oder regelgenerierendes Verfahren). Vgl. PhiloLex, <http://www.philoLex.de/indudedu.htm> (6.4.2020).

Turing-Test und Chinese Room Experiment

Wichtig für die Geschichte der KI war der Aufsatz „Computing Machinery und Intelligence“ (1950) des britischen Mathematikers A.M. Turing. „Intelligenz“ wird hier definiert als „die Reaktion eines intelligenten Wesens auf die ihm gestellten Fragen“.

Dieses Verhalten kann nach Turings Ansicht durch einen Test festgestellt werden (sog. Turing-Test): Eine Testperson kommuniziert über ein Computerterminal mit zwei ihr nicht sichtbaren Partnern, einem Programm und einem Menschen. Kann diese Person dabei nicht zwischen dem Menschen und dem Programm unterscheiden, kann das Programm als intelligent bezeichnet werden. Turings Ansatz wurde aus verschiedenen Gründen kritisiert. Ein wichtiger Kritikpunkt lautet, dass der Test nur auf Funktionalität, nicht jedoch auf Intentionalität und Bewusstsein prüft.

Die gegensätzliche Auffassung zu Turings Ansatz lautet, dass menschliche Intelligenz grundsätzlich nicht durch ein Computerprogramm ersetzt werden kann. Das versucht der Philosoph John S. Searle mit seinem Gedankenexperiment „Chinese Room Experiment“ zu verdeutlichen: Man soll sich vorstellen, eine Person befindet sich mit einigen Texten in chinesischer Sprache in einem geschlossenen Raum. Diese Person kann weder Chinesisch sprechen oder schreiben, noch ist sie in der Lage, die chinesischen Schriftzeichen als solche zu erkennen. Dieser Person werden nun durch einen Schlitz in der Wand Papierschneipsel mit einer Geschichte auf Chinesisch gereicht, dazu erhält sie Fragen in chinesischer Sprache zur Geschichte sowie ein „Handbuch“ in ihrer Muttersprache. Durch das Handbuch wird die Person in die Lage versetzt, anhand der erhaltenen Symbole (also der Geschichte und der Fragen) eine Antwort auf Chinesisch zu schreiben. Sie folgt dabei ausschließlich den Anweisungen aus der Anleitung und versteht somit die Antworten nicht, die sie wieder durch den Schlitz nach draußen schiebt. Dort steht ein chinesischer Muttersprachler, der die Geschichte und die Fragen formuliert hat. Nach den Antworten könnte er zu dem Schluss kommen, dass sich im Raum jemand befindet, der Chinesisch sprechen kann.

Searle will anhand dieses Gedankenexperiments zeigen, dass ein Programm, das den Turing-Test besteht, dadurch nicht zwangsläufig auch intelligent ist – es erscheint nur intelligent. Seit der Veröffentlichung des Chinesischen Zimmers gab es etliche Versuche, Searles Argument zu entkräften, auf einige der Gegenargumente geht Searle auch in seinem Werk „Minds, brains and programs“ (1980) ein.

Quellen: Klaus Manhart: Eine kleine Geschichte der Künstlichen Intelligenz, 17.1.2018, <https://www.computerwoche.de/a/eine-kleine-geschichte-der-kuenstlichen-intelligenz,33305372> (4.8.2019); Das Chinesische Zimmer, in: Künstliche Intelligenz. Informatik und Gesellschaft 2008/2009, http://www.informatik.uni-oldenburg.de/~iug08/ki/Grundlagen_Chinese_Room.html (13.8.2019).

Mustererkennung vs. Verstehen. Die Begrenztheit der Maschinenintelligenz veranschaulichte Schmid an einem Beispiel: Wenn man dem künstlichen Sprachassistenzsystem Siri (Software von Apple) folgende Anweisung gebe: „Sag mir was Schmutziges“, komme die Antwort: „Humus. Kompost. Schlamm. Schotter. Bimsstein.“ Die Antwort sei rein datenbasiert und auf Mustererkennung ausgerichtet und würde von einem Menschen so nicht gegeben, weil er den dahinterliegenden doppelten Wortsinn versteht.

Teilgebiete der KI. Bei allen Teilgebieten der KI (automatisches Schlussfolgern, Wissensrepräsentation, Planen, Sprachverarbeitung, maschinelles Lernen) zeigen sich nach Schmid immer zwei Probleme: Wie kann Wissen oder wie können Daten repräsentiert werden? Und wie können darin Antworten gesucht bzw. gefunden werden?²¹

²¹ Das weltweit am meisten in der Lehre genutzte Buch zu diesen Fragen ist von Russell Norvig: Artificial Intelligence (auch in Deutsch), Handbuch der KI (5. Aufl. 2020).

Zentrale Begriffe

Digitalisierung beschreibt einen Prozess, bei dem analoge Informationen (Texte, Objekte etc.) in digitale Formate umgewandelt werden, die dann mit Computerprogrammen verarbeitet werden. Davon profitiert auch die KI, insbesondere der datenintensive Teil (z.B. Deep Learning). Die KI kann umgekehrt auch die Digitalisierung unterstützen, z.B. durch Optical Character Recognition (OCR), die automatisierte Texterkennung innerhalb von Bildern.

Der Begriff **Big Data** kommt aus der Datenbankforschung. Hier geht es vor allem darum, Technologien zur Verarbeitung großer Datenmengen zu entwickeln und thesenfrei Muster in großen Datenmengen zu entdecken. KI, insbesondere maschinelles Lernen, liefert hilfreiche Methoden für die Analyse großer Datenmengen. Viele Daten bedeuten aber nicht immer automatisch bessere Ergebnisse: KI fokussiert im Kern nicht so sehr auf die Nutzung vieler Daten, sondern der passenden bzw. relevanten Daten für Schlussfolgerungen.

Bei **maschinellern Lernen/Mustererkennung** handelt es sich um einen Teil der KI. Die Lernverfahren aus beiden Traditionen sind größtenteils verschmolzen: Die Mustererkennung stammt aus der Signalverarbeitung, maschinelles Lernen aus der KI.

Data Science ist eine Interdisziplin aus Statistik, Datenbanken, Information Retrieval (Abrufen von komplexen Informationen) und maschinellem Lernen.

Robotik kommt überwiegend nicht aus der KI-Forschung, sondern aus dem Maschinenbau (in Bezug auf Industrieroboter). Teilgebiete der KI sind jedoch die Bereiche Kognitive Robotik und Mensch-Roboter-Interaktion.

Quelle: Vortrag Prof. Dr. Ute Schmid auf der Konferenz

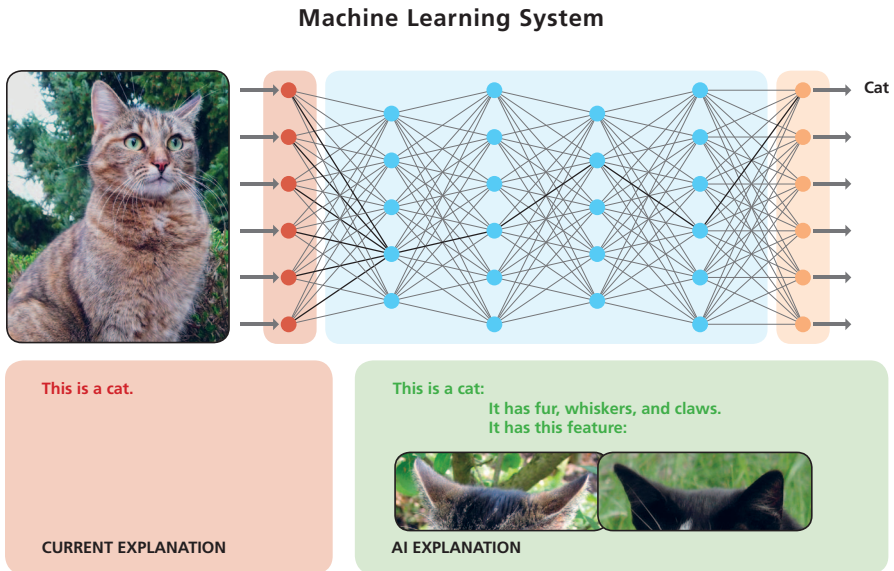
Nachvollziehbarkeit und Kontrolle. Nach Schmid ist es für das zukünftige Zusammenwirken von Menschen und KI wichtig, dass Entscheidungen der KI nachvollzogen und kontrolliert werden können, insbesondere wenn es sich um sicherheitskritische Bereiche handelt, z.B. in der Medizin oder in der Automobilindustrie.²² Dabei stelle sich jedoch die Frage, wer Entscheidungen der KI überhaupt nachvollziehen und kontrollieren kann, wenn diese selbstständig lernt und arbeitet. Folgende Erkenntnisse sollten nach Schmid dabei berücksichtigt werden:

- Viele Ansätze des maschinellen Lernens, insbesondere Tiefe Neuronale Netze sind Blackboxes.
- Beim maschinellen Lernen sind immer auch Biases²³ vorhanden, wie auch beim menschlichen Lernen, weil ohne Biases nicht gelernt werden kann.
- Sampling Biases sind nicht zu vermeiden (Diskriminierung aufgrund selektiver Datenlage). Wenn man Stichproben von Daten zieht und daraus lernt, weiß man häufig gar nicht, wie die darunterliegende echte Verteilung der Daten aussieht. Nimmt man z.B. immer nur Beispiele von Stoppschildern im Sonnenschein wahr, kann das Stoppschild bei Regen als solches nicht mehr erkannt werden.
- Gelerntes ist nie korrekt.
- Zentrale Anforderungen an Machine Learning-Systeme sind geringe Generalisierungsfehler (insbesondere geringe „false alarm rate“), Transparenz (z.B. für Sicherheitsbeurteilungen bei autonomen Systemen wie selbstfahrenden Autos), Nachvollziehbarkeit und Vertrauenswürdigkeit von Systementscheidungen.

²² Zudem gibt es regulatorische Anforderungen der EU-Datenschutzgrundverordnung (DSGVO), die nicht nur die komplette Kontrolle der Nutzer_innen über ihre (personenbezogenen) Daten sicherstellen sollen, sondern auch festlegen, dass (algorithmische) Entscheidungsprozesse durch künstliche Intelligenz transparent, nachvollziehbar und erklärbar sein müssen. Vgl. Datenschutzbeauftragter-Info: Das Date zwischen Datenschutz und der Künstlichen Intelligenz, 29.4.2019, <https://www.datenschutzbeauftragter-info.de/das-date-zwischen-datenschutz-und-der-kuenstlichen-intelligenz/> (6.4.2020).

²³ Bias bedeutet kognitive Verzerrung oder systematischer Fehler. Bei der Informationsverarbeitung entsteht ein Bias aus unzureichender Versuchsplanung/Datenlage oder durch Rückgriff auf Heuristiken (Urteilsfehler). Vgl. Bias. In: Lexikon der Psychologie, Spektrum.de, <https://www.spektrum.de/lexikon/psychologie/bias/2350> (6.4.2020).

Wellen von KI-Forschung. Schmid erläuterte, dass die Defense Advanced Research Projects Agency (DARPA)²⁴ drei Wellen von KI ausgerufen hat, aktuell die dritte: Die erste Welle war „Beschreiben“ (handgefertiges Wissen bzw. wissensbasierte Techniken), die zweite „Kategorisieren“ (statistisches Blackbox-Lernen) und die dritte steht am Beginn von „Erklären“ (kontextabhängige Adaption), das durch ein partnerschaftliches, interaktives und transparentes Vorgehen gekennzeichnet ist. Zur erklärbaren KI (Explainable Artificial Intelligence, XAI) gehören auch sogenannte Companion Systeme oder partnerschaftliche Systeme, bei denen Blackbox-maschinelles Lernen und wissensbasierte Ansätze kombiniert werden.



Quelle: <http://www.darpa.mil/program/explainable-artificial-intelligence>

24

DARPA (Defense Advanced Research Projects Agency) ist eine Behörde des Verteidigungsministeriums der Vereinigten Staaten, die Forschungsprojekte für US-Streitkräfte durchführt, <https://www.darpa.mil/about-us/darpa-perspective-on-ai> (6.4.2020).

Mensch-KI-Partnerschaft. Schmid forscht im Bereich human-level machine learning und favorisiert eine partnerschaftliche KI: „Wichtig bei KI ist kognitive Befähigung der Menschen und nicht kognitive Verdummung“, sagte Schmid.

„Wichtig bei KI ist die kognitive Befähigung der Menschen.“

Tabelle 1: Verhältnis Mensch – Maschine bei Entscheidungen

	Entscheidung liegt ...	
beim Menschen	gemeinsam bei Mensch und Maschine	bei der Maschine
Assistenzsystem	Partnerschaftliches System	Autonomes System

Quelle: Eigene Darstellung

Erklärbare künstliche Intelligenz. Partnerschaftliche Systeme ermöglichen die Generierung von Erklärungen: Hier wird z.B. erforscht, wie eine Assistenz zum Löschen oder Archivieren von irrelevanten Dateien entwickelt werden kann (DFG-Projekt, SPP Intentionales Vergessen), die über ihr Lösch- und Archivierverhalten Auskunft gibt. Grundlage ist ein interaktives Konzept: Es basiert darauf, das System fragen zu können, aus welchen Gründen es eine bestimmte Datei als irrelevant vorschlägt. Die Antwort könnte dann dafür genutzt werden, nicht überzeugende Teile der Begründung zu entfernen und so dafür zu sorgen, dass das System weiterlernt und verbessert wird. Partnerschaftliche Systeme könnten z.B. auch medizinisches Fachpersonal in Bezug auf Schmerzverstehen unterstützen, etwa durch kontrastive Beispiele (BMBF-Projekt TraMeExCo²⁵, DFG-Projekt PainFaceReader²⁶).

KI in Deutschland. Auch in Deutschland hat KI schon eine relativ lange Geschichte. So gibt es schon mehrere Jahrzehnte bei der Gesellschaft

25 Vgl. <https://www.uni-bamberg.de/en/cogsys/research/projects/bmbf-project-trameexco/> (6.4.2020).

26 Vgl. <https://www.uni-bamberg.de/en/cogsys/research/projects/dfg-projekt-painfacereader/> (6.4.2020).

für Informatik den Fachbereich Künstliche Intelligenz (FBKI), der auch die Zeitschrift Künstliche Intelligenz herausgibt.²⁷ Wolfgang Bibel von der Technischen Universität Darmstadt war einer der ersten KI-Wissenschaftler. Nach Ansicht von Schmid steht die KI-Forschung in Deutschland derzeit wissenschaftlich gut da, gemessen an der Zahl der wissenschaftlichen Publikationen zum Thema und der Zitationshäufigkeit. Allerdings wurden seit Anfang des Jahres 2000 kaum Lehrstühle mit der Bezeichnung „Künstliche Intelligenz“ besetzt. In dieser Zeit der „KI unter anderem Namen“ wurden entsprechend ausgerichtete Stellen als „Kognitive Systeme“ oder „Intelligente Systeme“ bezeichnet.

Politische Empfehlungen. „KI birgt große Chancen, aber auch große Risiken“, sagte Schmid. In Zeiten des Hype müsse man besonders achtsam sein: Dann sei Besonnenheit statt Aktionismus gefragt. Zum Plan der Bundesregierung, im Rahmen der KI-Strategie hundert KI-Professuren einzurichten, merkte sie an, dass es gegenwärtig nicht genügend Wissenschaftler_innen mit ausreichender Expertise in KI gebe. Wenn man jetzt zu schnell eine große Zahl an Professuren (fehl-)besetze, sei das Risiko sehr groß, dass auf diesen Stellen über Jahrzehnte schlechte Wissenschaftler_innen forschen und lehren. Viel besser wäre es, die KI-Expertise von unten aufzubauen, z.B. durch Doktorand_innenprogramme. Auch sollten Expert_innen aus der Wissenschaft in die Konzeption und Umsetzung entsprechender Programme eingebunden werden, insbesondere solche, die KI und Machine Learning auch in der Lehre vertreten. Die KI-Bildung sollte unbedingt auch in Schulen und in der Fortbildung für Unternehmen intensiviert werden.

Zunehmend sind Kooperationen von Universitäten mit großen Digitalplattformen festzustellen, wie z.B. der Universität Tübingen mit Amazon und der Technischen Universität München mit Facebook. Grundsätzlich findet Schmid Kooperationen von Universitäten und internationalen KI-Firmen durchaus sinnvoll, weil dadurch zusätzliche Kompetenzen und Mittel ins System kommen. Allerdings müssten dabei bestimmte Vorgaben eingehalten werden. Die Politik müsse in diesem Bereich stark regulieren. Daten, die Unternehmen

„Die KI-Forschung sollte stärker als bisher an Hochschulen gefördert werden.“

27 Der Fachbereich ist Träger der wissenschaftlichen Arbeit der Gesellschaft für Informatik auf allen Teilgebieten der Künstlichen Intelligenz (KI), wobei sowohl die theoretischen, die software- und hardwaretechnischen Aspekte und die Verbindungen zu Anwendungsgebieten abgedeckt werden. Vgl. <https://fb-ki.gi.de/> (6.4.2020).

von ihren Kunden erheben, sollten grundsätzlich nicht nur dem Unternehmen, sondern allen, insbesondere auch der Forschung, zur Verfügung stehen – natürlich mit Ausnahme unternehmenseigener, geheimnissensibler Daten. Auch wäre es wichtig, die KI-Forschung stärker als bisher an Hochschulen zu fördern. Dabei sollte man nicht nur anwendungsorientierte Forschung im Blick haben, sondern insbesondere auch Grundlagenforschung ohne direktes Ziel der Verwertbarkeit. Insgesamt müsse die Attraktivität der KI-Wissenschaft an Hochschulen erhöht werden, um einer Abwanderung von Doktorand_innen in große Firmen wie Google entgegenzuwirken. Dringend notwendig sei zudem der Aufbau einer funktionierenden Hardware-Infrastruktur im eigenen Land (Cluster, Clouds).

Diverse und breite Förderung als Ziel. Die Politik sollte bei ihrer Förderung darauf achten, KI divers zu fördern, also nicht nur auf Maschinelles Lernen/Deep Learning zu setzen. Gezielt sollten auch andere Bereiche gegen den Trend gestärkt werden, wie z.B. Wissensrepräsentation und automatisches Schlussfolgern. Auch sollten nicht nur wenige Exzellenz-Standorte mit KI ausgebaut, sondern in die Breite der Wissenschaftseinrichtungen investiert werden. Sonst drohe die Gefahr, dass KI-Bubbles entstehen und der durchschnittlich hohe Standard an Hochschulen in Deutschland gesenkt wird. Ein weiterer wichtiger Punkt der Förderpolitik sollte sein, andere relevante Informatikthemen nicht zu ignorieren, in denen Deutschland traditionell stark ist, etwa Sicherheit, Korrektheit, Algorithmik und Programmiersprachen.

KI sollte als Partner verstanden werden, um mehr Zeit für Menschlichkeit zu haben.

All diese Maßnahmen sind nach Schmid notwendig, um Deutschland im Bereich KI zukunftsfähig aufzustellen.

Soziale und demokratische KI. Schmid verdeutlichte neben den Potenzialen auch die Grenzen von KI: „KI löst keine gesellschaftlichen Probleme, kann aber helfen, Komplexität zu meistern und Menschen zu entlasten.“ Notwendig sei in jedem Fall eine soziale KI: KI sollte als Partner verstanden werden, um mehr Zeit für Menschlichkeit zu haben – und nicht, um Menschen bzw. Arbeitsplätze wegzurationalisieren. Schmid wandte sich gegen ein Menschenbild, das die Menschen als Datenlieferant_innen und manipulierbare Käufer_innen oder Wähler_innen betrachtet. Angestrebt werden sollte ein ganzheitliches Menschenbild, das den Menschen ins Zentrum der verschiedenen gesellschaftlichen Bereiche stellt, etwa in den Bereichen Bildung, Medizin oder Umwelt. Schmid setzt sich auch für eine demokratische KI ein: Wichtig wäre es, die Teilhabe von allen an den Daten und eine Entlastung durch KI zu erreichen.

BEISPIELE AUS DER ANWENDUNG

KI in der Industrie

Dr. Ulli Waltinger, Leiter der Forschungsgruppe Artificial Intelligence (AI) Lab bei Siemens²⁸, erläuterte in seinem Vortrag, wie künstliche Intelligenz verantwortungsvoll in industriellen Anwendungen eingesetzt werden kann. Bei Siemens werden KI-Systeme in erster Linie zur Optimierung und Effizienzsteigerung, aber auch zur Objekt-, Fehler- und Stresserkennung in industriellen Produktionsprozessen und Produkten verwendet.

Arbeitsdefinition von KI. „Intelligenz wird als Fähigkeit verstanden, Modelle zu bauen, die eine bestimmte Dynamik erklären“, sagte Waltinger. Künstliche Intelligenz bedeute, zielgerichtete Computer-Modelle zu bauen, die menschliche kognitive Prozesse abbilden können, z.B. Sehen, Hören, Gehen, aber auch die Ableitung von Argumentationen. Unter zielgerichtet sei zu verstehen, dass diese Systeme z.B. ein Objekt erkennen, Sprache übersetzen oder eine Aktion planen können. Die künstliche Dimension ergebe sich aus der Beteiligung von Computern. KI solle Probleme lösen, die bisher nur von Menschen gelöst werden konnten, und sich dabei selbst immer mehr verbessern.

Vielfältige KI-Felder. KI umfasse verschiedene Felder: Wahrnehmen, Lernen, Argumentieren und natürliche Sprachen. Als wichtige Anwendungsfelder nannte Waltinger die Computer-Vision, die Stärkung von Bildung, wahrscheinlichkeitsberechnete grafische Modelle, Entscheidungen und Pläne, die Algorithmische Spieltheorie, Wissensrepräsentation, Robotertechnik, Rede und Dialog sowie Objekt-/Spracherkennung.

28 Der Industriekonzern Siemens arbeitet bereits seit mehr als 20 Jahren mit künstlicher Intelligenz und hat auch schon KI-Systeme in zahlreiche Produktionsprozesse und Produkte eingebaut. 2017 wurde das KI-Labor „AI Lab“ eingerichtet, das auf die Entwicklung von Algorithmen und Geschäftsmodellen mit KI ausgerichtet ist. Vgl. Jonas Jansen: Wie Künstliche Intelligenz der Industrie hilft, 9.10.2018, <https://www.faz.net/aktuell/wirtschaft/code-n/wie-kuenstliche-intelligenz-der-industrie-hilft-15828933.html> (6.4.2020).

Datengetriebene Verfahren und Wissensrepräsentation. Im derzeitigen Hype-Cycle von KI herrsche die Wahrnehmung vor, dass KI zu 99 Prozent aus datengetriebenen Verfahren besteht: Meist werde KI mit maschinellem Lernen im Bereich des neuronalen Lernens gleichgesetzt, das es ermöglicht, aus großen Datenmengen eine Zielfunktion zu optimieren (Deep Learning = tiefes Lernen mit künstlichen neuronalen Netzen)²⁹. Dies sei aber nur ein relativ kleiner Teilbereich des maschinellen Lernens, der nur etwa 10 bis 20 Prozent auf dem breiten Feld von KI ausmacht. Im Bereich der KI existiere immer schon eine Konkurrenz zwischen Wissensrepräsentation einerseits und datengetriebenen Verfahren (wie z.B. dem maschinellen Lernen) andererseits. Vermutlich sei der gegenwärtige Hype von datengetriebenen Verfahren in der Forschung bald zuende und Wissensgraphen, Wissensrepräsentation und kognitive Aspekte würden wieder stärker in den Vordergrund rücken.

Fortschritte bei Machine Learning und Deep Learning. Waltinger berichtete, dass der größte Fortschritt der KI in den letzten Jahren auf den besonderen Fortschritten in den Bereichen maschinelles Lernen und Deep Learning basiert:

- Bei künstlicher Intelligenz sollen Maschinen dazu befähigt werden, Aufgaben so intelligent wie menschliche Wesen zu lösen.
- Beim maschinellen Lernen wird ein Set von Algorithmen so genutzt, dass intelligente Systeme aus Erfahrung lernen können.
- Beim Deep Learning werden Systeme aufgebaut, die tiefe neuronale Netze in einem großen Set von Daten so nutzen, dass sie nicht-lineare Veränderungen machen können.

Es habe etwa acht bis zehn Jahre gedauert, bis die Innovationen, die durch maschinelles Lernen und Deep Learning erzeugt wurden, in Industrieprozessen und Produkten ankamen, erläuterte Waltinger. In dieser Zeit habe sich auch die Datengrundlage signifikant verändert, d.h. es mussten erst genügend Daten vorhanden sein, um den Fortschritt zu ermöglichen. Neuronale Netze hätten zwar schon vorher existiert, doch konnten die Verfahren der Bild- und Spracherkennung durch eine verbesserte Datengrundlage erheblich verbessert werden. Durch Verfah-

29 Siehe Begriffsdefinition Deep Learning, Info-Kasten, Einführung in dieser Publikation.

ren des Deep Learning könne nun Wissen aus unstrukturierten Texten extrahiert werden, d.h. es ist möglich, die internen Prozesse und Produkte durch automatische Wissensgraphen-Extrahierung anzureichern.

Kooperation von Industrie und Universitäten. Nach Waltingers Eindruck sind sich Industrie und Universitäten auf dem Feld der KI-Forschung in letzter Zeit näher gekommen und kooperieren immer enger. Dadurch sei auch der Innovationszyklus zwischen dem Zeitpunkt der ersten Studie und der Umsetzung im Produkt deutlich kürzer geworden. In der Schnittmenge etablierten sich auch neue Formen der Zusammenarbeit, z.B. Entrepreneur Center. Die Unternehmen hätten erkannt, dass junge Talente aus den Universitäten großen Geschäftssinn und gute Kompetenzen mitbringen, die sie zügig und engagiert einsetzen möchten. Auf dem Feld der Kooperation von Unternehmen und Universitäten zeige sich eine unglaubliche Geschwindigkeit. Gepaart mit einer Open Source-Initiative (offene Quellcode-Initiative)³⁰ werde diese Entwicklung zu einem sehr starken Treiber, der signifikante Veränderungen anstoße.

Waltinger betonte, dass künstliche Intelligenz in der industriellen Anwendung mit erheblichen Potenzialen verbunden ist. Sie werde hier vor allem dazu genutzt, smarte³¹ Entscheidungen zu ermöglichen und die Performance und Präzision zu verbessern. Eine große Herausforderung bestehe darin, in industriellen KI-Anwendungen die gesamte Komplexität von Prozessen zu erfassen, vor allem die „unknown Unknowns“ (vorhandenes, aber unbekanntes Unbekanntes). Bei Siemens wird KI auf vielfältige Weise eingesetzt: vom selbstständigen Optimieren von Gasturbinen, verbessertem Monitoring von smarten Stromnetzen bis hin zu vorausschauender Wartung von industriellen Anlagen.

30 Open Source ist im allgemeinen Verständnis ein Paradigma, nach dem Wissen für alle offen und verwendbar ist. Im engeren Sinne werden damit Computerprogramme und -tools bezeichnet, deren Programmcode für jeden Interessierten frei einsehbar ist. Open-Source-Softwaresysteme können lizenzkostenfrei aus dem Web heruntergeladen und kostenfrei genutzt werden. Vgl. <https://www.arocom.de/fachbegriffe/open-source> (6.4.2020).

31 Die SMART-Formel im Kontext des Projektmanagements bedeutet: Spezifisch, Messbar, Attraktiv, Realistisch und Terminiert. Smarte Entscheidungen müssen demnach diese fünf Bedingungen erfüllen. Vgl. Heike Lorenz: Gute Entscheidungen treffen. In: Das Unternehmerhandbuch, 4.10.2010, <https://das-unternehmerhandbuch.de/gute-entscheidungen-treffen/> (6.4.2020).

Anwendungsbereiche von KI und Deep Learning bei Siemens – einige Beispiele

Produktion: Optimierung von Gasturbinen:

- Die NO_x (Stickoxide)-Emissionen von Gasturbinen können durch den Einsatz von KI um 15 bis 20 Prozent reduziert werden. Machine Learning-Algorithmen optimieren die Kontrollparameter automatisch und besser als menschliche Expert_innen. Dadurch optimieren sich die Gasturbinen selbstständig.

Pilot: Fehlerlokalisierung in intelligenten Stromnetzen (Smart Grids)

- Das Stromversorgungsnetz braucht intelligente Geräte, die zu einer gleichmäßigen Stromerzeugung beitragen können. Intelligente Stromnetze (Smart Grids) kombinieren Erzeugung, Speicherung und Verbrauch von Strom. Eine zentrale Steuerung stimmt sie aufeinander ab und gleicht Leistungsschwankungen – insbesondere durch erneuerbare Energien – im Netz aus. Durch verbesserte Algorithmen, Echtzeit-Streaming und die Interpretation von in Stromnetzen gemessenen Daten können Fehler aufgespürt werden. Machine Learning kann bei ca. 70.000 Fehlervorfällen in Stromnetzen eingesetzt werden.

Eingeschränkter KI-Anwendungsfall: Qualitätsverbesserung in der Industrie 4.0-Fabrik Amberg

- Es können mehr als 120 Varianten eines Produkts an einem Tag hergestellt werden, basierend auf 75 Prozent Automation. Im Ergebnis können weniger als elf Mängel auf eine Million (dpm) erreicht werden, was einer Qualitätsquote von 99,9988 Prozent entspricht.

Permanenter Ausbau: Erhöhung der Qualität und der Produktion in AM (Additive Manufaktur, 3-D-Drucken für Metall)³²

- Die Produktionsstunden im Additive Manufacturing (AM)-Betrieb können mit KI erhöht werden. Machine Learning-Algorithmen ent-

32

„Additive Fertigungsverfahren sind Fertigungsverfahren, bei denen das Bauteil – im Gegensatz zu subtraktiven Verfahren – durch Hinzufügen von Volumenelementen oder Schichten direkt aus digitalen 3D-Daten automatisiert aufgebaut wird oder auf einem bestehenden Werkstück wei-

decken und klassifizieren automatisch Themen. Der Vorteil besteht darin, dass die Produktionsstunden verdoppelt und gleichzeitig während des Druckens qualitative Probleme identifiziert werden können. Ein extra Zwischenschritt der Qualitätskontrolle fällt weg, da diese während der laufenden Produktion stattfindet.

Quelle: Vortrag Dr. Ulli Waltinger auf der Konferenz



Smart Grids (Foto: Sergey Nivens/Adobe Stock)

tere Volumenelemente aufgebaut werden. Wesentliches Merkmal aller Verfahren ist der Entfall produktspezifischer Werkzeuge und Vorbereitungen (werkzeuglose Fertigung).“ (Martin Kumke: Grundlagen der additiven Fertigung. Methodisches Konstruieren von additiv gefertigten Bauteilen, S. 7-23, In: AutoUni – Schriftenreihe book series (AUS, Bd. 124), 15.5.2018, https://link.springer.com/chapter/10.1007%2F978-3-658-22209-3_2 (6.4.2020).

Antrainieren von mehreren Aufgaben. Eine wichtige Frage ist derzeit nach Waltinger, wie erreicht werden kann, dass ein KI-System mehrere Aufgaben gleichzeitig trainiert. Bisher fokussieren sich Machine Learning- und Deep Learning-Systeme auf eine Zielfunktion (z.B. auf ein Objekt, eine Klassifikation, eine Entität). KI-Systeme, die mehrere Zielfunktionen kombinieren, können z.B. bei Ausschreibungssoftware eingesetzt werden, wenn Themen definiert und gezielt bestimmte Mitarbeiter_innen gefunden werden sollten. Unternehmen wollten solche Verfahren meist nicht mehr manuell durchführen, sondern die Maschine solle für sie herausfinden, welche Expert_innen für ein bestimmtes Thema richtig sind. Solche datenbasierten Verfahren seien geeignet, um große Datenmengen durchzuwühlen und Muster herauszukristallisieren, meinte Waltinger. Das könne nicht nur bei Ausschreibungen, sondern auch bei Texten und Angeboten sehr hilfreich sein.

Inzwischen sei es auch möglich, riesengroße Netze mit hunderten Millionen von Parametern zu bauen, wie z.B. Google und Facebook. Es sei jedoch schwierig, solche Netze wieder „klein“ zu bekommen bzw. High Performing-neuronale Netze so zu komprimieren, dass sie auf einem Controller oder Produkt liegen. Eine entscheidende Frage sei dabei, wie smarte Maschinenlernsysteme auf proprietäre³³ Systeme übertragen werden können, sodass sie verlässlich für die nächsten zwanzig Jahre halten. „Die Konnektivität ist für uns eine große Herausforderung“, sagte Waltinger. Siemens-Sensoren werden z.B. nach ihrer Installation bis zu dreißig Jahre genutzt. Deshalb sei es wichtig, lernende Algorithmen und die notwendige Rechenpower auch in Zukunft auf diese Geräte transportieren zu können.

Waltinger betonte, dass bei Siemens das Maschinenlernen, Deep Learning und Robotics in jedem einzelnen Bereich von der Produktentwicklung über die Produktion bis hin zu den Human Resources (Personalwesen) und im Servicebereich Einzug gehalten haben und auch immer wichtiger für das Unternehmen werden. Dabei gibt es nach Waltinger einige Herausforderungen, die das Unternehmen zusammen mit den Hochschulen erforscht.

33 Proprietäre Systeme (Closed Source) werden von freier Software (Open Source) abgegrenzt. Sie sind eigentums geschützt bzw. nicht frei und quelloffen. Im der Regel handelt es sich um eigene herstellereigenspezifische Entwicklungen, die nicht allgemeinen Standards entsprechen. Oft sind sie nicht oder nur mit Schwierigkeiten von Dritten implementierbar und deshalb auch nicht oder nur schwer zu öffnen oder zu lesen, weil sie z.B. lizenzrechtlich oder durch Patente beschränkt sind oder durch spezielles Know-how keinen Zugang gewähren. Vgl. Webcampus: Open Source vs. Proprietäre Systeme, <https://www.webcampus.de/blog/142/open-source-vs-proprietäre-systeme-was-ist-das-richtige-fuer-mich> (6.4.2020).

Forschungstrends und Herausforderungen. Waltinger benannte zentralen Trends in der KI-Forschung und skizzierte die damit verbundenen Herausforderungen.

1. KI und Datenknappheit: Wie sollen wir mit ungenügenden oder nicht gelabelten Daten umgehen?

Trends:

- Durch Übertragen lernen
- Von reichen Simulationen lernen
- Daten generieren (generative Modelle)

2. KI und Assistenzfunktionen (Companion) zur Entscheidungsunterstützung: Wie können Systeme korrekt eingesetzt werden, sodass sie den Mitarbeiter_innen am besten nützen und sie bei verschiedenen Aufgaben unterstützen? Wann sollte ein System etwas vorschlagen und entscheiden, wann sollte auf jeden Fall der Mensch hinzugezogen werden, um den Vorschlag zu validieren und den nächsten Entscheidungsschritt zu machen?

Trends:

- Modelle von Menschen und ihren Aufgaben entwickeln
- Modell der Komplementarität von Mensch und Maschine entwerfen
- Initiativen von Mensch und KI-Systemen koordinieren

3. KI und Gesellschaft: Wie kann erklärbare KI erreicht werden? Was macht ein System und aus welchen Gründen?

Trends:

- Vertrauen und Sicherheit schaffen
- Gerechtigkeit und Transparenz sicherstellen
- Ethische und rechtliche Unabhängigkeit wahren
- Arbeitsplätze und Wirtschaft berücksichtigen

Erklärbarkeit hat nach Waltinger eine große Bedeutung, wenn in Industrieprozessen komplexe Entscheidungen getroffen werden müssen. Es sei sehr wichtig zu wissen, warum ein System oder neuronales Netz eine bestimmte Entscheidung trifft. „Erklärbarkeit ist die Basis für einen verantwortungsvollen Umgang mit KI und betrifft verschiedene Ebenen: die Algorithmen selbst, den Ort der Installation, aber auch die komplexe Systemebene“, sagte Waltinger.

„Erklärbarkeit ist die Basis für einen verantwortungsvollen Umgang mit KI.“

4. KI und Offene Welt: Wenn Training und Echtwelt zwei verschiedene Bereiche sind – wie können verlässliche Voraussagen von „unknown Unknowns“ erreicht werden?

Trends:

- Erweiterte Echtwelt-Tests durchführen
- Algorithmische Portfolios erstellen
- Ausfallsichere Designs schaffen
- Menschen und Maschinen koordinieren

Verantwortungsvoller Umgang mit KI. Waltinger betonte, dass ein verantwortungsvoller Umgang mit Technologie und Echtwelt ein absolutes Alleinstellungsmerkmal von Europa und Deutschland sein sollte. Im Siemens Lab werden im Rahmen der Initiative „Verantwortlicher Umgang mit KI“ Best Practice-Richtlinien erarbeitet, um einen verantwortungsvollen Umgang mit KI im Unternehmen zu gewährleisten. Darüber hinaus müssen Unternehmen ethische Richtlinien umsetzen, die eine Expert_innen-gruppe im Auftrag der EU-Kommission erarbeitet hat.

Ethische Richtlinien für den Einsatz von KI in Europa: Ansätze und Kritik

Die EU-Kommission beauftragte im Sommer 2018 die „Hochrangige Expertengruppe zur KI“ (AI HLG), Richtlinien für eine vertrauenswürdige KI zu erarbeiten. Am 8. April 2019 legten die 52 Kommissionsmitglieder das Ergebnis vor. Zweck der Leitlinien ist die Förderung einer vertrauenswürdigen KI bei der Anwendung und Entwicklung in Europa. Die Leitlinien beginnen mit der Feststellung, dass eine vertrauenswürdige KI drei Komponenten umfasst, die während des gesamten Lebenszyklus des KI-Systems gegeben sein sollten. Sie muss

- rechtmäßig sein und alle geltenden Gesetze und Vorschriften einhalten,
- ethisch sein und die Einhaltung ethischer Grundsätze und Werte gewährleisten,
- sowohl aus technischer als auch aus sozialer Sicht robust sein, da KI-Systeme auch bei guten Absichten unbeabsichtigte Schäden verursachen können.

Grundsätze sind

- die Anerkennung der Handlungsautonomie des Einzelnen (Achtung der menschlichen Selbstbestimmung),
- das Prinzip „do no harm“ (Vermeidung von Schaden),
- Fairness,
- Nachvollziehbarkeit/Erklärbarkeit.

Benannt werden sieben Schlüsselanforderungen, die KI-Systeme erfüllen sollten, um als vertrauenswürdige gelten zu können:

1. Menschliches Handeln und Kontrolle: KI-Systeme sollen den Menschen zu fundierten Entscheidungen befähigen und seine Grundrechte fördern.
2. Technische Robustheit und Sicherheit: KI-Systeme müssen mit einem präventiven Ansatz gegenüber Risiken entwickelt werden.
3. Privatsphäre und Datenqualitätsmanagement: KI-Systeme müssen

die Privatsphäre und den Datenschutz während des gesamten Lebenszyklus eines Systems gewährleisten.

4. Transparenz: Die Geschäftsmodelle für Daten, Systeme und KI sollten transparent sein.

5. Vielfalt, Nichtdiskriminierung und Fairness: Unlautere Verzerrungen müssen vermieden werden.

6. Gesellschaftliches und ökologisches Wohlergehen: Die breitere Gesellschaft und die Umwelt sollten als Interessenträger während des gesamten Lebenszyklus des KI-Systems betrachtet werden.

7. Rechenschaftspflicht: Es müssen Mechanismen geschaffen werden, um Verantwortung und Rechenschaftspflicht für KI-Systeme und ihre Ergebnisse zu gewährleisten.

Am 19. Februar 2020 stellte die EU-Kommission in ihrem „Weißbuch zur künstlichen Intelligenz“ ihre Strategie für ein digitales Europa vor, das bis zum 19. Mai 2020 zur öffentlichen Konsultation steht.

Einschätzung und Kritik

Nach der Veröffentlichung der Leitlinien kritisierte ein Mitglied der Expertengruppe, Prof. Dr. Thomas Metzinger, das Ergebnis scharf. Der Professor für theoretische Philosophie an der Universität Mainz war der Auffassung, dass die Leitlinien zwar derzeit das Beste seien, was es weltweit zum Thema verantwortungsvoller Umgang mit KI gibt, doch handle es sich um einen Kompromiss, der viel zu industriefreundlich, zu kurzfristig und zu vage sei. Die Geschichte von einer „verantwortungsvollen KI“ sei ein von der Industrie erdachtes Marketing-Narrativ. In Wirklichkeit gehe es darum, Zukunftsmärkte zu entwickeln und Ethikdebatten als „elegante öffentliche Dekoration für eine groß angelegte Investitionsstrategie“ zu benutzen. Der Leitgedanke einer „vertrauenswürdigen KI“ sei begrifflicher Unsinn, da Maschinen nicht vertrauenswürdig sein könnten, sondern nur Menschen. Nicht-vertrauenswürdige Konzerne oder Regierungen könnten sich weiterhin unethisch verhalten, auch wenn sie eine gute, robuste KI-Technologie besitzen.

Als Teil des Problems kennzeichnete Metzinger die Zusammensetzung der Expert_innengruppe, die überwiegend aus Vertreter_innen der Industrie und nur vier Ethiker_innen bestand. Entsetzt zeigte sich Metzinger darüber, dass das EU-Expertengremium nicht dazu bereit war, echte rote Linien zu ziehen, also nicht verhandelbare ethische Prinzipien für die Verwendung von KI in Europa festzulegen, etwa ein Verbot von tödlichen autonomen Waffensystemen, der KI-gestützten Bewertung von Bürger_innen durch den Staat (Social Scoring) im Stil Chinas, aber auch grundsätzlich eines Einsatzes von KI, die von den Menschen nicht mehr verstanden und kontrolliert werden könne.

Nach Metzinger muss eine gute KI auf jeden Fall eine ethische KI sein. Deshalb bestehe nun eine moralische Verpflichtung, die Richtlinien selbst aktiv weiterzuentwickeln, die – bei aller Kritik an ihrer Entstehung und dem Ergebnis – gegenwärtig die beste Grundlage für die nächste Phase der Diskussion seien. Ihre juristische Verankerung in europäischen Grundwerten sei ausgezeichnet und die erste Auswahl der abstrakten ethischen Prinzipien zumindest passabel.

Nun sei es höchste Zeit, dass sich die Universitäten und die Zivilgesellschaft intensiv an diesem Prozess beteiligen und der Industrie die Hoheit über den Ethik-Diskurs wieder aus der Hand nehmen. Man befinde sich in einem rasanten historischen Übergang, der auf vielen Ebenen gleichzeitig stattfindet. Das Zeitfenster, in dem dieser Wandel zumindest teilweise noch zu kontrollieren sei und die philosophisch-ethischen Grundlagen der europäischen Kultur wirksam verteidigt werden könnten, werde sich in wenigen Jahren schließen.

Quellen: High-Level Expert Group on Artificial Intelligence, Ethics Guidelines for Trustworthy AI, 8. April 2019, <https://ec.europa.eu/futurium/en/ai-alliance-consultation>; Europäische Kommission: Weißbuch zur künstlichen Intelligenz, Brüssel, 19.2.2020, COM(2020) 65 final, https://ec.europa.eu/info/sites/info/files/commission-white-paper-artificial-intelligence-feb2020_de.pdf; Thomas Metzinger: Nehmt der Industrie die Ethik weg!, In: Der Tagesspiegel, 8.4.2019, <https://www.tagesspiegel.de/politik/eu-ethikrichtlinien-fuer-kuenstliche-intelligenz-nehmt-der-industrie-die-ethik-weg/24195388.html>; Monika Ermert: Ethische Grenzen für künstliche Intelligenz. In: Heise Magazine, 10/2019, <https://www.heise.de/select/ct/2019/10/1557133582448874> (20.08.2019).

Beispiele für den Einsatz von KI-Systemen. Waltinger nannte Beispiele, mit welchen Möglichkeiten und Problemen der Einsatz von KI-Systemen für Menschen und Unternehmen verbunden sind:

- **Autonome technische Systeme:** In Berlin wurde die erste autonome Tram als Forschungsprojekt präsentiert. Kollisionswarnsysteme und smarte Assistenten helfen, die Sicherheit zu erhöhen und die Nutzbarkeit zu verbessern. Aber wie kann gewährleistet werden, dass Maschinen genauso Entscheidungen treffen wie es ethisch handelnde Menschen tun würden? Wenn man z.B. eine autonome Tram in einer Stadt auf die Reise schicken will, müsse man sich sicher sein, dass die ganze Komplexität des Systems berücksichtigt wird und zugleich Sicherheit gewährleistet ist. Inwieweit soll das System autonom entscheiden können und welche Rolle sollen Entscheidungen von Menschen spielen?



Autonom fahrende Metro in Alicante (Spanien) (Foto: Wikipedia)

- **Kenntnisreicher virtueller Assistent,** z.B. in der Bilddiagnostik: Radiolog_innen können bei zweifelhaften Fällen durch KI-basierte Bildanalyse dabei unterstützt werden, mit wachsenden diagnostischen Erfordernissen umzugehen, z.B. bei der Entdeckung von Krankheiten mit (fast) 100-prozentiger Genauigkeit. Aber wie soll darauf reagiert werden, dass sich die Rollen und Jobprofile rasch wandeln und wem können und wollen die Patient_innen ihre Daten anvertrauen?

- Großflächiges Monitoring in Gebäuden, etwa zur Nutzungsart oder Wohndauer: Das Wissen, wie viele Menschen zu einem bestimmten Zeitpunkt in einem Raum oder einem Gebäude sind, ist hoch relevant, weil dadurch z.B. der Nutzerkomfort erhöht und die Energiekosten signifikant reduziert werden können. So kann anhand dieser Kriterien entschieden werden, ob gelüftet oder geputzt werden sollte. Aber welche Information muss dafür wann weitergeleitet werden? Wie viele und welche Menschen müssen informiert werden? Ein Problem sei dabei, dass privatsensitive Daten (v.a. Videos) notwendig sein könnten, um die erforderlichen Informationen zu erhalten. Doch sei es notwendig, die Privatsphäre des Einzelnen zu wahren bzw. die Anforderungen des Datenschutzes zu erfüllen.

Mit all den Möglichkeiten von KI seien somit auch Gefahren des Missbrauchs und Risiken verbunden, z.B. die Kontrolle von Personen und Räumen oder Profiling auf der Ebene von Produkten, Personen und Unternehmen. Hier müssten Gegenmaßnahmen ergriffen werden, um negative Entwicklungen zu vermeiden, und es müsse auch transparent gemacht werden, wie ein Geschäftsmodell erfolgreich und zugleich verantwortungsvoll im jeweiligen Kontext umgesetzt werden soll.

Nutzungsmöglichkeiten. Waltinger verdeutlichte, dass KI-Lösungen mit vielfältigen Nutzungsmöglichkeiten verbunden sind, z.B.

- visuelle Überwachung,
- Profiling von Individuen, Produkten und Unternehmen,
- autonome System- und Verfahrensoptimierung,
- Generierung und Kreativität,
- autonome physische Kontrolle,
- digitale Begleitung bei System- und Verfahrensoptimierung,
- Wissensbeschaffung und -vorschläge.

Gewinne. Daraus könnten sich zahlreiche Gewinne ergeben, z.B.

- neue Niveaus von Innovation und Kunst,
- höhere Systemeffizienz,
- leichteres Komplexitätsmanagement,
- mehr Transparenz und Sicherheit,
- mehr Prozesseffizienz,
- höhere Geschwindigkeit und Qualität der Arbeit,
- Befreiung der Menschen von wiederholenden Tätigkeiten,
- Reduktion von Kosten und Fehlern,
- schadensfreie Infrastruktur.

Risiken. Es seien aber auch zahlreiche Risiken damit verbunden, z.B.

- widerrechtliches Eindringen in Bereiche und Manipulation,
- Black Box/Intransparenz,
- Missbrauch von persönlichen Daten,
- Risiken von Verzerrungen und Diskriminierung,
- Irrtümer und Sicherheitsrisiken,
- Verlagerung/Abschaffung von Arbeitskräften,
- Unfähigkeit der Kontrolle von KI,
- utilitaristische Entscheidungen (rein nutzenorientiert),
- Verstärkung von Ungleichheiten.

Ziele der Regulierung. Die Technologie selbst, aber auch Regulierung könnten dabei helfen, Risiken zu minimieren und die richtige Balance zu finden. Wichtige Ziele wären:

- Systemsicherheit
- Ziel- und Wertevereinbarungen
- Kontrollmechanismen
- Datenschutz/Schutz der Privatsphäre
- Erklärbarkeit und Transparenz
- Verantwortlichkeit und Verpflichtung
- Geteilter Benefit (Nutzen, Vorteile)

Interdisziplinäre, multifunktionale Teams. Bei Forschung und Entwicklung im KI-Bereich ist nach Waltinger Co-Creation von entscheidender Bedeutung. Dabei handelt es sich um einen Prozess, der von oben – der Führungsebene – stimuliert wird, dann beschäftigen sich – auf der Arbeitsebene – interdisziplinäre Teams mit dem Thema. Die Mitglieder der Arbeitsgruppen kommen aus unterschiedlichen Bereichen und Disziplinen, z.B. dem Machine Learning, der Softwareentwicklung oder dem Kundenbereich. Nur mit einem interdisziplinären cross-funktionalen Team könnten in einem Unternehmen End-to-End-Prozesse³⁴ gelingen, meinte Waltinger.

34

„End-to-End-Prozesse sind funktions- und bereichsübergreifende Geschäftsvorfälle vom Kunden zum Kunden durchgängig durch das Unternehmen, verbunden mit klaren Definitionen, mit dem Ziel die gesamte Wertschöpfungskette effektiv und effizient zu gestalten und zu steuern. End-to-End Prozesse ermöglichen das Heben von Potenzialen in den Schnittstellen funktionaler Grenzen.“ Vgl. Ralf Gaydoul/Christian Daxböck: Prozessmanagement von End-to-End-Prozessen, In: Controlling & Management review, Sonderheft 2/2011, Springer Professional, <https://www.springerprofessional.de/prozessmanagement-von-end-to-end-prozessen/6403860> (6.4.2020).

Mitarbeiter_innen im Bereich KI. Für ein Unternehmen sei es zwar relativ leicht, qualifizierte Mitarbeiter_innen mit der richtigen Hochschulausbildung für KI zu gewinnen, wie z.B. im Bereich Maschine Learning. Doch sei es schwierig, diese Talente längerfristig an ein Unternehmen zu binden. Junge Mitarbeiter_innen wollten in der Regel an „echten Weltproblemen“ arbeiten und sich nicht auf Themen wie „Katzenklassifikation“ beschränken. Deshalb sei es vorteilhaft, wenn Unternehmen in Projekten konkrete Probleme mit gesellschaftlicher Relevanz definieren, den Talenten aus den Hochschulen Freiräume geben und sie mit anderen Kolleg_innen in einem interdisziplinären und agilen Team zusammenbringen. Die große Herausforderung liege darin, die passenden cross-funktionalen Teams zusammenzustellen und sie langfristig zu halten.

Aus Sicht von Waltinger ist die Universitätsausbildung mittlerweile sehr gut. Die große Zahl der Absolvent_innen bringe sehr gute Kompetenzen, enormen Esprit und auch Engagement mit. Die Entrepreneurship-Zentren würden den Absolvent_innen ein großes Mindset mitgeben, z.B. Wissen über Design Thinking oder Business Valuation. Allerdings bestehe in Deutschland das Problem, dass es – im Unterschied zu den USA – keine große Venturing-Szene gibt, die junge Menschen in ihren innovativen Vorhaben unterstützt.

Umgang mit Daten und Gefährdungen. Waltinger verwies auf die Charta of Trust der Industrie, in der festgehalten ist, wie im Industriesektor miteinander kollaboriert werden sollte und was unter Sicherheit (auch Cybersicherheit) zu verstehen ist. Schon jetzt sei Profiling im Consumer-Bereich sehr wichtig: Die meisten Menschen hätten hier schon eine Nummer, doch wüssten sie nicht, wie die Algorithmen ausgerichtet sind und was z.B. passiert, wenn sie eine Versicherung abschließen möchten. Algorithmen würden bereits über vieles entscheiden – und es sei davon auszugehen, dass diese Tendenz immer mehr zunimmt. „Daten sind ein enormes Asset“, sagte Waltinger.

Nach Waltinger birgt KI aber auch die Chance, die notwendigen Daten zu wesentlichen Problemen einer breiten Community zur Verfügung zu stellen. „Das ist eine Herausforderung und eine Möglichkeit: Man kann mit KI die richtigen Probleme gemeinsam auf Basis der richtigen Daten bearbeiten“, sagte Waltinger.

„Man kann mit KI die richtigen Probleme gemeinsam auf Basis der richtigen Daten bearbeiten.“

KI in der Medizin

Über den Stand der Forschung und die Möglichkeiten der KI in der Medizin berichtete Prof. Dr.-Ing. Anja Hennemuth, Charité – Universitätsmedizin Berlin, Institut für kardiovaskuläre Computer-assistierte Medizin. Sie arbeitet interdisziplinär an der Schnittstelle zwischen der Charité, der Technischen Universität Berlin und dem Fraunhofer Institut für digitale Medizin.

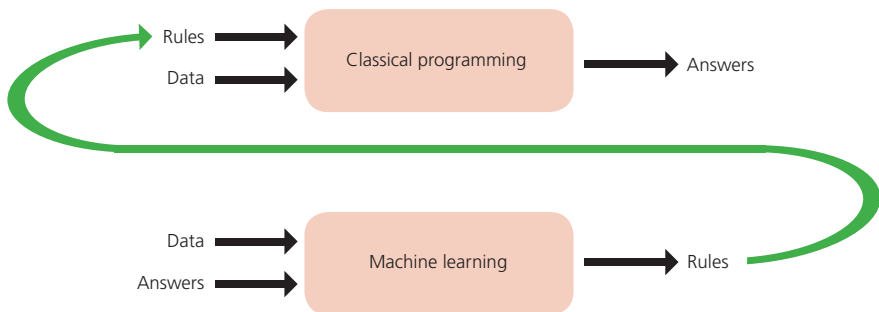
Diskrepanz zwischen Forschung und Praxis. Auch in der Medizin sei inzwischen die gestiegene Bedeutung der KI angekommen, meinte Hennemuth. Das zeige sich z.B. darin, dass jeden Monat neue KI-Journals aus dem Boden schießen. Allerdings sei eine große Diskrepanz zwischen dem Hype in der Forschung und den Entwicklungen in der Praxis festzustellen, in der die Digitalisierung oft noch auf Skepsis stoße. Laut einer Studie würden Ärzt_innen nach wie vor sehr ungern mit Computern arbeiten, was vor allem auf Gefühle der Depersonalisierung zurückzuführen sei. Diese Zögerlichkeit von Ärzt_innen im Umgang mit Computern zeige sich auch an der (mangelnden) Qualität oder Unvollständigkeit der verfügbaren Daten, die als Basis für das maschinelle Lernen oder KI-Systeme gebraucht werden.

Hennemuth betonte, dass Computer und Digitalisierung die Qualität in der Medizin erheblich steigern können. Maschinen seien in manchen Bereichen besser als Menschen, z.B. beim Bildvergleich, der beim Hautkrebs-Screening oder bei Radiografiebildern zur Identifikation von Lungentumoren eingesetzt wird. Entsprechend komme KI in der Medizin vor allem in diesen Feldern zum Einsatz: bei der Automatisierung von Prozessen, die für Ärzt_innen anstrengend, schwierig und fehleranfällig sind. Methoden des Maschinellen Lernens seien auch bereits in einer Vielzahl klinisch eingesetzter Produkte verfügbar.

Klassische Anwendungen von schwacher KI. Die klassischen KI-Lösungen bestehen in der Extraktion relevanter Bildeigenschaften und dem Trainieren von Klassifikatoren. Ein wichtiges Feld ist nach Hennemuth die Computer-unterstützte Entdeckung von Tumorherden. Schon Anfang der 2000er Jahre erfolgte die erste Zulassung von Produkten, die frühe Formen des einfachen maschinellen Lernens mit statistischen Klassifikatoren darstellten. Damit können auffällige Strukturen, z.B. Lungenknoten, Mamma-Kalzifikationen oder Darmpolypen in digitalen Bilddaten entdeckt, markiert und klassifiziert werden.

Klinische Entscheidungsunterstützung. Zu den klassischen Ansätzen der KI-Anwendung in der Medizin gehören zum einen wissensbasierte Systeme, zum anderen physiologische und biophysikalische Modelle. Im Bereich wissensbasierter Systeme werden klinische Leitlinien in Software implementiert. So ist es möglich, eine große Menge an Regelwerken schnell zu durchforsten bzw. Regeln zur Diagnostik aus Büchern zu extrahieren, um Ärzt_innen in Entscheidungsprozessen zu unterstützen. Bei physiologischen und biophysikalischen Modellen geht es um die mathematische Implementierung physikalischer Gesetze. Dies kann z.B. dazu genutzt werden, in Studien die Wirkung von Medikamenten vorherzusagen, therapeutische Interventionen vorher durchzuspielen oder zu erfahren, wie sich die Physiologie von Patient_innen bei körperlicher Anstrengung verändern würde.

Was ist anders mit Maschinellern Lernen?



Quelle: Eigene Darstellung in Anlehnung an: François Chollet: Deep Learning with Python, Shelter Island 2018, S. 5

Maschinelles Lernen. Ein entscheidender Unterschied zwischen klassischem Programmieren und maschinellern Lernen ist, dass die Regeln an anderer Stelle stehen. Wenn z.B. in der Medizin noch kein Wissen über die Regeln in einem bestimmten Bereich existiert und Prädiktionen (Vorausagen) über Verläufe gemacht werden sollen, kann maschinelles Lernen dazu beitragen, die Regeln zu finden, die Ärzt_innen dann interpretieren können. In Deutschland werden auch häufig Simulationen eingesetzt. Diese können bei unvollständigen Daten und Messungen sehr hilfreich sein, indem Patient_innen unterschiedliche Untersuchungs- und Behandlungsabfolgen durchlaufen.

Hennemuth erläuterte, wie maschinelles Lernen im Sinne von Deep Learning funktioniert: Zunächst werden Daten gesammelt, bereinigt und organisiert. Dann wird ein Ziel festgelegt, ein passendes Modell und eine geeignete Trainingsstrategie ausgewählt. Das Modell wird an das angepasst, was in den Daten gesehen wird. Schließlich wird das Modell trainiert – mit einer Evaluation der Modell-Komplexität und der Modell-Performance. So kann das optimale Modell für die jeweiligen Daten herausgearbeitet werden.

Von außen betrachtet sehe es häufig so aus, als wäre in der Medizin die Brücke in die Zukunft schon geschlagen bzw. künstliche Intelligenz flächendeckend einsatzbereit, meinte Hennemuth. In der Praxis werde aber erst in kleinen Schritten auf dieses Ziel hingearbeitet und es müssten noch viele Bausteine zusammengesetzt werden. KI in der Medizin sei in vielen Fällen nur eine erweiterte Mustererkennung und die damit verbundenen großen Potenziale würden bei Weitem noch nicht genutzt.

Bildbasierte Diagnostik. Hennemuth nannte ein verbreitetes Anwendungsbeispiel in der bildbasierten Diagnostik: Maschinelles Lernen kann z.B. in der Kardiologie dazu genutzt werden, genauer zu erkennen, was im Ultraschall zu sehen ist, wodurch deutliche Verbesserungen in der Diagnostik erreicht werden können. Die vollautomatische Analyse und Interpretation von Ultraschalldaten folgt einem bestimmten Ablauf: 1. Orientierungserkennung (in der Ansicht auf das Herz), 2. Segmentierung (Sichtbarmachen der Strukturen des Herzens), 3. Vermessung der Strukturen, 4. Erkennen von Erkrankungen.

Ein anderes Feld, in dem maschinelles Lernen unterstützend wirkt, ist die Echtzeit-Bildgebung. Eine neue Bildgebungstechnologie (MRT als Echtzeit-Messung) ermöglicht die Darstellung von Bewegung in hoher zeitlicher Auflösung, z.B. des Herzens in verschiedener Anstrengung. Eine manuelle Auswertung wäre unmöglich, da bis zu 900 Bildzeitpunkte in 20 Schichten in verschiedenen Anstrengungsstadien angezeigt werden. Die KI hilft dabei, die zahlreichen medizinischen Bilddaten zu analysieren und die diagnostischen Parameter auszusuchen. Weitere Einsatzmöglichkeiten gibt es z.B. in der Radiologie.

Verzichtbarkeit von Ärzt_innen? Immer wieder taucht die Frage auf, ob Ärzt_innen durch KI-Systeme verzichtbar werden. Während die einen meinen, dass das künftig sehr gut möglich sein könnte, widersprechen andere deutlich: Maschinen würden niemals fähig sein, die interrelationale Qualität der therapeutischen Beziehung von Ärzt_innen und Pa-

tient_innen herzustellen. „Der persönliche Kontakt zwischen Ärzt_innen und Patient_in ist so wichtig und einzigartig, dass KI ihn nicht übernehmen kann“, sagte Henemuth. Die Unterstützung durch KI könnte bei bestimmten Aufgaben aber dazu führen, dass Ärzt_innen dadurch mehr Zeit haben, um sich mit Patient_innen intensiver zu beschäftigen und längere Gespräche zu führen.

Unterstützung in Notaufnahmen und auf Intensivstationen. Weitere wichtige Einsatzfelder für KI in der Medizin sind Notaufnahmen und Intensivstationen in Krankenhäusern, wo das Personal unter Zeitdruck zahlreiche Untersuchungen durchführen muss. An beiden Orten besteht das Problem, dass Ärzt_innen viele Tätigkeiten ausführen müssen, für die sie keine Spezialist_innen sind. Hier erweist sich der Einsatz von KI bzw. maschinellem Lernen als äußerst sinnvoll: KI-Systeme können bei der Interpretation von Daten unterstützen, chirurgische Roboter durch Automatisierung wiederholte zeitkritische Aufgaben übernehmen. Deshalb betrachtet Hennemuth KI auf der Intensivstation als besonders wertvoll: An diesem Ort sei die Arbeitsbelastung sehr hoch, Multitasking notwendig und die Arbeitssituation von vielen Sensorsystemen und Fehlalarmen geprägt. „Durch KI-Unterstützung kann nicht nur reaktiv, sondern auch proaktiv gehandelt werden, bevor die Komplikation da ist“, sagte Hennemuth. Diese Art der Anwendung sei besonders erfolgversprechend: Die Komplexität der Information könne an die KI ausgelagert werden, während sich die Ärzt_innen der Rettung der Patient_innen zuwenden können.

„Durch KI-Unterstützung kann nicht nur reaktiv, sondern auch proaktiv gehandelt werden.“

Früherkennung von Komplikationen. KI bietet auch eine wichtige Unterstützung bei der Früherkennung: Indem neuronale Netze mit Routinedaten trainiert werden, können Komplikationen (z.B. Blutungen, Nierenversagen) schon zu einem frühen Zeitpunkt vorhergesagt werden. So kann rechtzeitig gegengesteuert und auch der Tod von Patient_innen verhindert werden. Die Ergebnisse durch KI-Systeme seien deutlich besser als der aktuelle Standard, meinte Hennemuth.

KI in der Forschung. In der Forschung werden aufgrund neuer Messverfahren (z.B. der Proteomik³⁵) oder verbesserten Bildgebungsmethoden

³⁵ Ein Proteom umfasst die Gesamtheit aller Proteine, die in einer Zelle oder einem Lebewesen unter definierten Bedingungen vorliegen. Die Proteomik oder Proteomanalyse umfasst verschiedene Verfahren, um ein Proteom mit biochemischen Methoden systematisch zu einem bestimmten

heute viel mehr Informationen als früher erfasst, die dabei helfen können, Krankheiten zu spezifizieren und zu charakterisieren. Damit sei auch die Hoffnung verbunden, durch eine verbesserte Phänotypisierung eine bessere Unterteilung von Patient_innen mit unterschiedlichen Bedürfnissen zu erreichen. Das sei wichtig für pharmazeutische Entwicklungen und eine geeignete bzw. gezieltere Therapieauswahl, aber auch bei der Individualisierung von medizinischer Diagnose und Therapie.

Zulassung von Medizin-Produkten. Die Zulassung von KI-basierten Medizinprodukten werde oft als schwierig angesehen. Zu unterscheiden ist hier zwischen Systemen, bei denen die Entscheidung weiterhin beim Arzt bzw. bei der Ärztin liegt und die Software lediglich ein unterstützendes System bereitstellt (mittleres Risiko), und Systemen, die die Gesundheit direkt beeinflussen, wie z.B. bei einem Herzschrittmacher (hohes Risiko). Die Bausteine bei der Zulassung sehen vor, dass ein Qualitätsmanagementsystem (mit Standardprozessen für Entwicklung, Risikomanagement und Testing), eine technische Dokumentation (Manual, Designspezifikation, Systemarchitektur) sowie eine klinische Evaluation durchgeführt wird. Die Prozessunterstützungstools seien beim maschinellen Lernen hervorragend, meinte Hennemuth. Ein wesentliches Problem sei jedoch, dass eine KI immer wieder mit aktuellen Daten trainiert werden muss. Das bedeute, dass auch die für die Zulassung wichtige klinische Evaluation immer wiederholt werden muss.

Die europäische Datenschutz-Grundverordnung (DSGVO) und der Standard ISO 27001 geben vor, dass immer die Möglichkeit bestehen muss, Ergebnisse zurückzuverfolgen, die durch ein Produkt geliefert werden.³⁶ Der KI werde hier häufig vorgeworfen, sie wäre eine Black Box, meinte Hennemuth, was aber nicht der Fall sei. Für den Menschen sei es nur relativ schwierig, innerhalb tiefer neuronaler Netze die Datenrepräsentation zu verstehen. Daraus ergebe sich eines der wichtigsten Forschungsgebiete: „Wir brauchen eine erklärbare und regulierte KI“, sagte Hennemuth.

Zeitpunkt zu analysieren. Vgl. Proteomik. In: Lexikon der Biologie, Spektrum der Wissenschaft, <https://www.spektrum.de/lexikon/biologie/proteomik/54197> (6.4.2020).

36 Am 25. Mai 2018 ist die neue Datenschutz-Grundverordnung (DSGVO) der EU in Kraft getreten. Sie verpflichtet Unternehmen, die notwendigen technischen und organisatorischen Maßnahmen zu ergreifen, um ein hohes Maß an Informationssicherheit gemäß Artikel 32 zu gewährleisten: Sicherheit der Verarbeitungsdaten. ISO 27001 ist der internationale Standard für Informationssicherheit und beschreibt die Best-Practice-Anforderungen für die Implementierung eines Informationssicherheitsmanagementsystems (ISMS). Vgl. Wie ISO 27001 Ihnen hilft, Ihre Daten zu schützen. Vgl. <https://www.itgovernance.eu/de-de/gdpr-compliance-with-iso-27001-de> (6.4.2020).

Erklärbare und regulierte KI. Bei einer erklärbaren KI geht es darum, dass die Aktivitäten und Verarbeitungsschritte eines Netzes verständlich werden, indem die Informationen benannt werden, die zur Entscheidung geführt haben: Welche Informationen waren für das Ergebnis des Netzes relevant? Nach Hennemuth ist diese Frage in der Medizin besonders wichtig, und es werde auch schon im Rahmen der Zentren für maschinelles Lernen (z.B. in München und Berlin) zum Thema erklärbare KI geforscht. Bei der regulierten KI wird versucht, schon bekannte Regeln und Gesetze in tiefe neuronale Netze einzubringen, um sicherzustellen, dass die Sicherheitsbestrebungen erfolgreich sind und keine Ergebnisse entstehen, die völlig unplausibel sind oder in unerwartete Risiken führen. In der Medizin ist das schon vielfach Thema der Forschung, z.B. in der Histologie und Neurologie.

„Wir brauchen eine erklärbare und regulierte KI.“

KI in Medizinprodukten. Hennemuth berichtete, dass die Entwicklungs- und Zulassungsbedingungen von KI in Medizinprodukten international sehr unterschiedlich sind. In den USA können z.B. nach dem 21st Century Cures Act von 2016 Softwareprodukte anders als Geräte betrachtet werden und die US-Arzneimittelbehörde FDA unterstützt die Zulassung KI-basierter Produkte. In der EU soll ab Mai 2020 auch Software, die Vorschläge berechnet, als medizinisches Gerät betrachtet werden.

Auch das Teilen von Daten, die eine wichtige Basis für KI darstellen, sei aufgrund der Datenschutzbestimmungen in Deutschland und Europa nicht einfach. Für eine zukunftsfähige Ausbildung von Studierenden, die sich für das Machine Learning in der Medizin meist sehr interessieren, wäre es aus Sicht von Hennemuth sehr wichtig, dass sie früher die Möglichkeit haben, mit Echtweltdaten in Kontakt zu kommen bzw. zu arbeiten. Sehr lobenswert seien einzelne Initiativen, z.B. Challenges³⁷, bei denen alle Hürden der DSGVO überwunden werden sollen, um repräsentative klinische Daten der Allgemeinheit zugänglich zu machen.

Maschinelles Lernen auf klinischen Daten. Prospektive Studien mit personenbezogenen Daten/Proben oder biometrische Studien am Menschen sind alle bei der Ethikkommission beratungspflichtig. Bei Studien mit der Nutzung retrospektiv gewonnener Daten/Proben gibt es verschiedene

³⁷ Z.B. eine Übersicht über die Grand Challenges (großen Herausforderungen), die im Bereich der medizinischen Bildanalyse organisiert wurden, <https://grand-challenge.org/challenges/>, <https://www.kaggle.com/competitions> (6.4.2020).

Verfahren, auch in Bezug auf die Notwendigkeit von Patienteninformation (PI) und Einwilligungserklärung (EK).

Maschinelles Lernen – Klinische Implementierung. In KI-Systemen gibt es auch den Ansatz, dezentrale Wissensmodule so zu konzipieren, dass sie potenziell den kompletten Datenstrom größerer Kliniken filtern und analysieren können, um daraus höherwertige Wissensmodule abzuleiten. Die Rohdaten verbleiben dabei in der Klinik. Der entstehende Wissensmittelpunkt (Knowledge Hub) steht dabei im Zentrum der Entwicklungen durch Bündelung und Validierung der dezentral generierten Wissensmodule. Wissensmodule und anwendungsspezifische „Apps“ werden den teilnehmenden Kliniken und Projekten über den Knowledge Hub dann dynamisch zur Verfügung gestellt. Die Datenverarbeitungs- und Speicherkonzepte berücksichtigen dabei streng die Bedingungen des Datenschutzes und der Datensicherheit.

Medizinische Datenverfügbarkeit für KI. Mit Unterstützung des US-amerikanischen National Institut of Health (NIH) ermöglicht Physionet den Zugriff auf wissenschaftliche Datensammlungen (Routinedaten, aber auch der Intensivstation, EKG etc.). Dazu gehört die Bereitstellung von Trainings- und Validierungsdaten als Challenges und der Vergleich von Lösungen. Die Datensätze sind meist eher klein.

„Das Datenproblem in der Medizin ist bei KI sehr groß“, sagte Hennemuth. Einige, teilweise auch EU-geförderte Initiativen beschäftigen sich mit der

„Das Datenproblem in der Medizin ist bei KI sehr groß.“

Frage, wie medizinische Daten DSGVO-konform, auch kryptografisch oder über Blockchain-Technologie zur Verfügung gestellt werden können – allerdings mit der Möglichkeit der Dateneigner_innen, diese jederzeit wieder zurückziehen zu können. Wenn man diese Daten im Kontext von maschinellem Lernen und KI nutzen möchte, sei es eine extrem große Herausforderung, diesen Prozess zu organisieren. Schließlich müsse jederzeit damit gerechnet werden, dass die Informationen und Prozesse, die einem Algorithmus oder Modell antrainiert wurden, immer wieder an die geänderten Wünsche von Dateneigner_innen angepasst werden müssen.

Bei der Organisation von medizinischen Daten und der Art und Weise, wie diese zum Wohl der Patient_innen genutzt werden können, besteht nach Hennemuth in Deutschland noch sehr viel Spielraum. Gegenwärtig sei die Tendenz zu beobachten, dass über Apps Massen von Daten – auch

Gesundheitsdaten – gesammelt werden, die Menschen freiwillig zur Verfügung stellen. Da hier nur sehr wenige Regeln und Kontrolle existieren, entstehe als Folge sehr viel „Datenschmutz“, aus dem viele falsche Informationen gezogen werden können. Auch ethische und rechtliche Probleme seien damit verbunden.

Im Bereich der medizinischen Forschung wäre mehr Beratungsunterstützung sinnvoll. Die verbreitete große Unsicherheit in Bezug auf Datenschutz verhindere auch eine gute Zusammenarbeit zwischen den Institutionen. Viele Mediziner_innen wüssten z.B. nicht, was in einer Patientenzustimmung stehen muss, damit Daten nicht nur für eine konkrete Studie, sondern auch später im Austausch mit anderen Forscher_innen genutzt werden können. Die Lähmungserscheinungen in diesem Bereich seien extrem groß. Positiv seien die Aktivitäten der Gesundheitszentren zu bewerten, die vom Bundesministerium für Bildung und Forschung gefördert werden, da sie sich unter anderem mit der Frage von Forschungsdaten und dem Umgang damit – auch in Bezug auf den Datenschutz und die Qualitätssicherung – intensiv beschäftigen. Solche Vorhaben sollten unbedingt verstetigt werden.

Notwendige Überzeugungsarbeit. Nach Hennemuth besteht eine wichtige Aufgabe darin, die Ärzt_innen davon zu überzeugen, dass KI Arbeit reduzieren bzw. arbeitsentlastend sein kann und das Wissen und die Präsenz der Ärzt_innen weiterhin unverzichtbar bleibt. Wissenschaftliche Forschung in Deutschland werde abgehängt, wenn sich die Mediziner_innen diesen neuen Wegen nicht öffnen. Allerdings sei diese Überzeugungsarbeit ein sehr langsamer Prozess und hänge stark vom Engagement Einzelner ab, z.B. dem Leiter einer Chirurgie, der diese Aktivitäten fördert und seine Mitarbeiter_innen motiviert, zu diesen Themen zu forschen und zu publizieren. Dann würden Ärzt_innen durchaus eine Begeisterung für KI-Systeme und die damit verbundenen Möglichkeiten entwickeln. Grundsätzlich sei es in festgefahrenen Strukturen immer schwierig, grundlegende Veränderungen umzusetzen – auch unabhängig von KI. Doch sei es aufgrund der zahlreichen Vorteile lohnend, sich für den KI-Einsatz in der Medizin einzusetzen.

Innovation. Da KI-Systeme auf Daten basieren, die die Abbildung vergangener Daten sind, stellt sich die Frage, wie in solchen Systemen Innovation stattfinden kann. Nach Hennemuth sind tiefe neuronale Netze so aufgebaut, dass nachtrainiert werden kann. Das vorherige Wissen werde zwar genutzt, doch könne es in KI-Systemen immer an neue Entwicklungen und Erkenntnisse angepasst werden.

Eine weitere Frage ist, wie vermieden werden kann, dass KI-Systeme Fehler produzieren, etwa falsch-positive Befunde. Hennemuth meinte, dass Fehler natürlich auch bei KI-Systemen nicht auszuschließen sind. Doch sei auch bei Menschen noch nie eine Eliminierung falsch-positiver Befunde erreicht worden. Entscheidend sei jedoch, dass KI-Systeme weniger falsch-positive Befunde erzeugen als ein Mensch. Man könne nicht verlangen, dass eine Maschine perfekte Ergebnisse liefert, wenn die Alternative sehr viel schlechter ist, meinte Hennemuth. Sie findet es bedenklich, dass Innovationen abgelehnt werden, weil sie nicht perfekt sind, obwohl sie gegenüber dem State of the Art einen Riesenfortschritt darstellen. Hier sollten sich die Menschen öffnen und den Fortschritt sehen – und nicht das Unperfekte bzw. die Mängel in den Vordergrund stellen. „Fehler von Maschinen stoßen oft auf weniger Toleranz als Fehler von Menschen,“ sagte Hennemuth. Das führe bei KI-Forscher_innen häufig zu Frust. Durch eine solche Haltung würden in der Medizin wertvolle Potenziale verschenkt und Innovationen verhindert.

KI als Kommunikationspartner

Prof. Dr. Philipp Cimiano leitet am Exzellenzcluster Kognitive Interaktionstechnologien der Universität Bielefeld die Forschungsgruppe Semantische Datenbanken. Cimiano möchte eine neue Perspektive auf KI vorstellen und den Blick auf dieses Thema weiten. Der derzeitige Diskurs zur KI sei sehr eingeschränkt, da im Fokus der Betrachtung meist die Automatisierung stehe, die den Menschen quasi ersetzen soll. KI könne aber auch zum Sparringspartner des Menschen werden, der den Menschen nicht einfach nachmacht, sondern ihn herausfordert, ihm Grenzen aufzeigt und ihm dabei hilft, sich weiterzuentwickeln und „besser“ zu werden. Die daraus resultierenden Möglichkeiten könnten für den Standort Deutschland sehr spannend sein.

Unterstützung bei komplexen Entscheidungsprozessen. Cimiano verdeutlichte, dass Entscheidungsprozesse in der gegenwärtigen Welt immer mehr zunehmen und zugleich komplexer werden, da immer mehr Faktoren und Aspekte dabei berücksichtigt werden müssen. Die Menschen stünden vor der Herausforderung, mit einer enormen, stark wachsenden Daten- und Informationsflut umzugehen. Gleichzeitig verstärkten sich komplizierte und globale Abhängigkeiten, die Entscheidungsprozesse sehr komplex machen. So seien bei Entscheidungen z.B. die verschiedenen Per-



Künstliche Neuronale Netze (Foto: metamorworks/Adobe Stock)

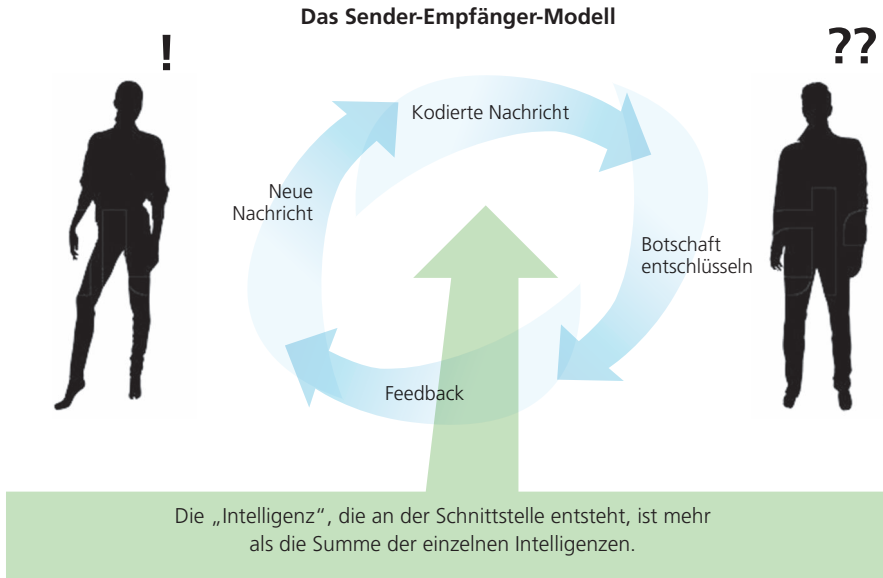
spektiven von Stakeholdern zu berücksichtigen, und es sei notwendig, zwischen relevanten, richtigen und falschen Informationen zu unterscheiden. Wichtig sei deshalb die Frage, wie die KI die Menschen bei der Entscheidungsfindung unterstützen kann: „KI soll den Menschen nicht ersetzen, sondern ihn dabei unterstützen, zu besseren Entscheidungen zu kommen“, sagte Cimiano.

„KI soll den Menschen nicht ersetzen, sondern ihn dabei unterstützen, zu besseren Entscheidungen zu kommen.“

Dabei sollte KI nicht zu einem Deus ex Machina werden, der bei einem Problem plötzlich herangezogen wird und eine Entscheidung treffen soll, die der Mensch dann nur noch bestätigt. Wesentlich sei vielmehr, wie erreicht werden kann, dass KI Entscheidungsprozesse transparent macht und die Menschen dahin führt, auf informierter Basis bessere Entscheidungen treffen zu können. Die dabei eingesetzten KI-Systeme müssten vom Menschen als kompetente Partner wahrgenommen werden.

Um die Interaktion zwischen Menschen kommunikationstheoretisch zu beschreiben, gibt es ein gängiges Kommunikationsmodell: das Sender-Empfänger-Modell.

INTERAKTIONSINTELLIGENZ



Quelle: Eigene Darstellung in Anlehnung an De Jaegher, H./Di Paolo, E./Gallagher, S. (2010) (siehe Fußnote 38).

Interaktionsintelligenz. Aus der Forschung ist bekannt, dass Kommunikation nicht nur reiner Informationsaustausch ist, sondern dabei ein „Mehr“ entsteht. In einer Definition der Kognitionswissenschaften entsteht Intelligenz an der Schnittstelle zwischen zwei rationalen Agenten, die miteinander interagieren. Die an der Schnittstelle entstehende „Intelligenz“ ist mehr als die Summe der einzelnen Intelligenzen: die Interaktionsintelligenz.³⁸ Aber wie kann die KI zu einem Partner werden, sodass an der Schnittstelle bei der Interaktion Mensch und Maschine Neues bzw. eine höhere Intelligenz entsteht?

Bisher können Maschinen so trainiert werden, dass sie gewisse Muster erkennen. Wenn eine Maschine z.B. mit Katzenbildern gespeist und somit

38

Vgl. De Jaegher H./Di Paolo, E./Gallagher, S. (2010): Can social interaction constitute social cognition? In: Trends in cognitive sciences, 14(10), S. 441-447.

trainiert wird, kann sie bei einem neuen Bild erkennen, ob es sich um eine Katze handelt. Dieses Vorgehen ist nach Cimiano aber viel zu restriktiv, um wirkliche Interaktionsintelligenz zu erzeugen, wie sie in der Kommunikation mit einem kompetenten Partner entstehen könnte. *„Das Ziel ist es, Menschen über KI zu befähigen.“*

Kompetente Partner. Nach Cimiano zeichnen sich KI-Systeme dann als kompetente Partner dadurch aus, dass sie

- ihre Position rechtfertigen und erklären können,
- tiefes Wissen über Anwendungsdomänen haben,
- ein Verständnis von Zusammenhängen und Wissen darüber haben, wie den Menschen Zusammenhänge erklärt werden können,
- die Perspektive des Gegenübers einnehmen können,
- die Simulation verschiedener Alternativen/Optionen durchführen können (What-if-Analysen),
- eine robuste und gewinnbringende Interaktion mit einem Menschen aufrechterhalten bzw. einen Dialog führen können,
- den Menschen bei Planung und Bewertung unterstützen können,
- fähig zur Quantifikation und Erklärung von Unsicherheit sind bzw. die Quelle der Unsicherheit benennen können.

„Das Ziel ist es, Menschen über KI zu befähigen und nicht, menschliche Unmündigkeit zu produzieren“, sagte Cimiano.

Forschungsprojekt Ratio. Cimiano leitet auch das DFG-finanzierte Forschungsprojekt RATIO (Robust Argumentation Machines). Das Schwerpunktprogramm strebt einen Paradigmenwechsel an, indem statt einzelner Fakten kohärente argumentative Strukturen die Informationseinheit bilden, die für die Entscheidungsfindung systematisch und explizit aufbereitet werden.³⁹ Die zentrale Frage lautet, wie Maschinen robust dazu gebracht werden können, Zusammenhänge einzuordnen und sie so aufzu-

39 Vgl. <http://www.spp-ratio.de/de/home/> (6.4.2020).

bereiten, dass sie menschliche Deliberationsprozesse⁴⁰ unterstützen können. Zahlreiche Disziplinen und Forschungsbereiche sind an dem Projekt beteiligt: Künstliche Intelligenz, Informationssuche, Wissensrepräsentation, Semantik, Computerlinguistik, Mensch-Maschine-Interaktion. Eine Grundannahme ist, dass mit diesem neuen Ansatz weitreichende Anwendungspotenziale im Bereich des Ingenieurwesens, der Medizin, des Finanzwesens und Online-Handels, der Politik sowie der Geistes- und Rechtswissenschaften erschlossen werden können. Das bundesweite Programm ist an zwölf Standorten bzw. Universitäten in Deutschland vertreten.

Wie können Maschinen dazu gebracht werden, so zu argumentieren, dass sie auf Augenhöhe mit Menschen interagieren können? Cimiano erläuterte, dass es dazu neuer Methoden bedarf, um die Argumente und ihre Zusammenhänge aus Dokumenten extrahieren zu können. Zudem müssen neue semantische Modelle und Ontologien zur tiefen Repräsentation von Argumenten entwickelt werden. Es werden neue Suchverfahren benötigt, die Argumente indexieren, einen öffentlichen Diskurs auswerten können bzw. die für eine Suchanfrage relevanten Für- und Gegenargumente finden sowie der gezielten Interaktion mit menschlichen Nutzer_innen zugänglich machen können, um deren Meinungsbildung zu unterstützen. Zudem sind neue Verfahren des maschinellen Schlussfolgerns zu entwickeln, um Implikationen von Argumenten und deren Plausibilität bewerten zu können. Dazu gehört die Synthese von Argumenten (interaktive Denkwerkzeuge) sowie die Bewertung von Argumenten (Erkennen von schlechter Argumentation, falschen Annahmen, Validierung von Fakten bzw. Identifizieren von Fake News/Alternativen Fakten).

Rationalisierung von Empfehlungen. Innerhalb des Programms RATIO koordiniert Cimiano das Projekt RecomRatio (Rationalizing Recommendations)⁴¹. Hier wird das Ziel verfolgt, eine computergestützte Methode zu entwickeln, die Menschen dabei unterstützen kann, eine Empfehlung systematisch gegen Alternativen abzuwägen und automatisch eine argumentative Synthese zu erstellen, nach denen die empfohlene Alternative als optimal zu sehen ist (sog. Empfehlungsrationalisierung). Dieses Verfahren könnte

40 Deliberation beschreibt eine auf den Austausch von Argumenten angelegte Form der Entscheidungsfindung unter Gleichberechtigten. Das bessere Argument soll zu besseren Entschlüssen führen, weil im Idealfall alle Argumente gegeneinander abgewogen werden und eine Einigung auf die „beste“ Lösung möglich ist. Vgl. Martin Große Hüttmann: Deliberation. In: Martin Große Hüttmann/Hans-Georg Wehling (Hrsg.): Das Europalexikon, 2., aktual. Aufl. Bonn: Dietz 2013, Bundeszentrale für politische Bildung, <https://www.bpb.de/nachschlagen/lexika/das-europalexikon/176777/deliberation> (6.4.2020).

41 Vgl. <http://www.spp-ratio.de/de/projekte/ratiorec/> (6.4.2020).

in der Medizin, aber auch in anderen Bereichen gewinnbringend eingesetzt werden. So wäre es möglich, dass das System in großem Stil durch Techniken des Machine Readings Tausende von Publikationen liest – z.B. zu klinischen Studien – und aus diesen Texten Argumente herausdestilliert, die eine bestimmte Therapieentscheidung untermauern können. Cimiano verdeutlichte, dass solch ein System einem Menschen, der eine Entscheidung treffen muss, als kritischer Partner gegenüberstehen und ihn bei einer systematischen Analyse des Entscheidungsraums unterstützen könnte.

Das Verfahren skizzierte er an einem Beispiel aus dem medizinischen Bereich:

- **Grundannahme:** Bei Menschen, die an der Augenerkrankung Glaukom erkrankt sind, ist eine Reduktion des täglichen mittleren Augeninnendrucks (IOP) wünschenswert.
- **Prämisse:** In elf vergleichbaren klinischen Studien hat sich gezeigt, dass die Behandlung mit 0,005 Prozent Latanoprost den täglichen mittleren IOP stärker senkt als eine Behandlung mit Timolol von 0,5 Prozent.
- **Schlussfolgerung:** Die Behandlung mit 0,005 Prozent Latanoprost ist effektiver als die Behandlung mit Timolol, um den täglichen mittleren IOP zu reduzieren.

Das System kann somit synthetisierte Argumente auf der Basis großer Datenmengen erzeugen. Damit soll erreicht werden, dass Entscheidungsträger_innen durch die Interaktion mit solchen Argumenten einen fundierteren Meinungsbildungsprozess durchlaufen können. Wichtig ist dabei nach Cimiano, dass die Argumente nicht statisch sind, sondern z.B. Mediziner_innen die Möglichkeit haben, bei einzelnen Patient_innen die gegenläufigen Abhängigkeiten zu verschieben: So können verschiedene Parameter zur Entscheidungsfindung herangezogen und unterschiedlich gewichtet werden. Wenn ein Nutzer z.B. der Sicherheit einen höheren Stellenwert geben möchte als der Wirksamkeit, können die Argumente systematisch verändert und dabei beobachtet werden, ob und wie sich die Schlussfolgerungen dadurch verändern. Ein solches KI-System könnte den Menschen somit dabei unterstützen, den Entscheidungsraum argumentativ zu durchdringen.

Cimiano betonte, dass Interaktionsintelligenz auf einer nicht-reduktionistischen Sichtweise auf KI basiert, d.h. KI nicht nur als Machine Learning

betrachtet. Dem Menschen werde nicht einfach eine Entscheidung abgenommen, sondern

- es können Zusammenhänge erfasst und bewertet werden (Kontext, Hintergrundwissen),
- es können Handlungsoptionen und Alternativen vorgeschlagen und eingeschätzt werden,
- es ist eine vorausschauende Planung möglich und
- es kann in Alternativen und Szenarien gedacht werden ("Was-Wenn"-Analysen).

Ergänzung für Wissenschaft. Dieses Verfahren sei dennoch keine Konkurrenz für die bisherige Wissenschaft, sondern könne sie sinnvoll ergänzen. In den meisten wissenschaftlichen Themenbereichen sei ein einzelner Mensch aufgrund unzähliger Studien weltweit gar nicht mehr in der Lage, die ganze Datenlage zu überblicken. Die notwendigen Fakten und Erkenntnisse seien zwar vorhanden, dem einzelnen Individuum aber nicht mehr zugänglich. Wissenschaftler_innen könnten trotz größter Bemühungen den Forschungsstand nicht mehr überblicken und wüssten auch nicht, welche Informationen wirklich vertrauenswürdig sind oder wie Biases erkannt werden können. Deshalb seien Entscheidungen sehr schwer geworden und es sei umso wichtiger, mit KI das individuelle Empowering in diese Richtung voranzutreiben. KI könne dabei helfen, die Faktenlage in einem Fach- oder Therapiegebiet systematisch zu ordnen und als Gesamtheit zu betrachten. KI als Sparringspartner könnte Forscher_innen z.B. darauf aufmerksam machen, dass bei einer Entscheidung noch nicht alle wichtigen Erkenntnisse beachtet wurden oder auf Studien mit gegenteiligen Ergebnissen verweisen. Die systematische Durch-

„Ein großer Unterschied zwischen Mensch und Maschine ist, dass Maschinen nicht kreativ werden.“

forstung der Studien lässt nach Cimiano ein Wissen entstehen, das eine gesicherte Entscheidungsbasis schafft als nur ein paar (zufällig ausgewählte) Studien, die der Mensch kennen und auswerten kann.

Selbstständiges Handeln von KI-Systemen? Muss eine KI, die einen Menschen inhaltlich unterstützen und argumentieren kann, nicht auch fähig sein, mit der Welt zu interagieren und über eigene Erfahrungen verfügen, um selbstständig Weltwissen aufzubauen? Wäre es prinzipiell denkbar, dass eine KI selber den Nobelpreis gewinnen könnte?

Cimiano beschäftigt sich hauptsächlich mit der KI-unterstützten Analyse unstrukturierter Daten, d.h. mit der Frage, wie aus textuellen Daten Erkenntnisse gezogen werden können, um Entscheidungen zu unterstützen. Dass eine Maschine den Nobelpreis gewinnen könnte, hält er für unrealistisch: „Ein großer Unterschied zwischen Mensch und Maschine ist, dass Maschinen nicht kreativ werden.“ Dies sei für eine exzellente Forschungsleistung aber unverzichtbar. Die Erwartungen wären zu hoch, wenn man solche Spitzenleistungen der Forschung in Maschinen hineinprojizieren würde. In seiner Forschung stehe die Frage im Mittelpunkt, wie in der Interaktion zwischen Mensch und Maschine Innovation entstehen kann. KI-Systeme könnten zwar neue Phänomene entdecken und dem Menschen Vorschläge machen, doch nehme ein Mensch die Vorschläge auf, bewerte sie und finde in Experimenten heraus, ob es der richtige Weg sein könnte. Der Mensch gebe der Maschine eine Richtung, und die Maschine könne dann über Kombinatorik bestimmte Dinge herausfinden. *„Entscheidend bleibt aber nach wie vor der Mensch.“*, sagte Cimiano.

KI-getriebene Veränderungen im wissenschaftlichen Publizieren

Im April 2019 hat der Verlag Springer Nature sein erstes maschinengeneriertes Buch veröffentlicht: „Beta Writer: Lithium-Ion Batteries. A Machine-Generated Summary of Current Research. Springer: Cham 2019.“ (1) Damit wurde Neuland betreten, weil zuvor noch kein maschinengeneriertes Buch von einem Wissenschaftsverlag veröffentlicht wurde. Der neue Buchprototyp im Fachbereich Chemie bietet einen Überblick über die neuesten Forschungspublikationen zum Thema Lithium-Ionen-Batterien. Das Ergebnis ist eine strukturierte, automatisch generierte Zusammenfassung einer großen Anzahl aktueller Forschungsartikel aus diesem Bereich (mehr als 50.000 Fachaufsätze zum Thema). Forscher_innen erhalten dadurch die Möglichkeit, das schnell wachsende Informationsaufkommen auf diesem Gebiet effizient zu überschauen. (2)

Das Projekt basiert auf einer Zusammenarbeit zwischen Springer Nature und Wissenschaftler_innen der Goethe-Universität Frankfurt/Main unter Leitung von Juniorprofessor Dr. Christian Chiarcos in der Angewandten Computerlinguistik. Die Wissenschaftler_innen entwickelten einen Algorithmus: Der sogenannte Beta Writer selektiert und verarbeitet automatisch relevante Publikationen, die auf der Plattform SpringerLink veröffentlicht wurden. Diese wissenschaftlich begutachteten Veröffentlichungen von Springer Nature werden von dem Algorithmus einem ähnlichkeitsbasierten Clustering unterzogen, um die Quelldokumente in zusammenhängende Kapitel und Abschnitte zu gliedern. Dadurch entstehen Zusammenfassungen auf Grundlage der publizierten Artikel. Die extrahierten Zitierungen sind mit Hyperlinks versehen, sodass die Leser_innen eindeutige Verweise auf die Quelldokumente erhalten. Automatisch erstellte Inhaltsverzeichnisse und Referenzen erleichtern die Orientierung innerhalb des Buchprototypen. (3)

Henning Schoenenberger von Springer Nature erwartet, dass im wissenschaftlichen Publizieren künftig verschiedene Inhaltstypen eine Rolle spielen werden: Das Spektrum reiche von vollständig von Menschen erstellten Inhalten über verschiedene Blended-Man-Machine-Textgenerierungen hin zu vollständig maschinengenerierten

Texten. (2) Nach Auffassung von Prof. Dr. Chiarcos dient diese Form maschinengenerierter Bücher hauptsächlich als Werkzeug, um menschliche Autor_innen dabei zu unterstützen, Texte zu schreiben. Ein automatisch generierter Literaturüberblick könne dazu genutzt werden, die (teilweise überbordenden und unübersichtlichen) Publikationen in einer bestimmten Wissensdomäne strukturiert zu erfassen und den aktuellen Forschungsstand abzubilden. (4)

Angesichts einer großen Flut an wissenschaftlicher Literatur – jedes Jahr werden ca. 2,5 Millionen wissenschaftliche Fachaufsätze publiziert – wird es für Forscher_innen zunehmend schwierig, in ihrem Fachgebiet die relevanten Publikationen herauszufiltern. KI-gestützte Systeme haben die Möglichkeit, große Textmengen durchzupflügen, die einzelne Forscher_innen nicht selbst sehen oder lesen könnten. Manche sehen deshalb in der automatisierten Erkenntnisproduktion eine große Chance für die Wissenschaft. KI-Systeme könnten der Forschung an der entscheidenden Stelle zum Durchbruch verhelfen. (5) So hat z.B. ein neuronales Netzwerk kürzlich Zusammenhänge in alten Fachaufätzen entdeckt, die den menschlichen Forscher_innen vorher nicht aufgefallen waren. (6)

Die Entwicklung geht weiter: Im April 2019 haben Wissenschaftler_innen des MIT (Massachusetts Institute of Technology) das Konzept eines neuronalen Netzwerks vorgestellt, das Fachaufsätze scannt und selbstständig einfache englische Zusammenfassungen erstellt – und auf diese Weise die Forscher_innen entlasten soll. (7)

Quellen: (1) Das Buch ist als kostenloser Download verfügbar: <https://link.springer.com/content/pdf/10.1007%2F978-3-030-16800-1.pdf> (2) Goethe-Uni online: Erstes maschinengeneriertes Buch publiziert, 4. April 2019, <https://aktuelles.uni-frankfurt.de/forschung/erstes-maschinengeneriertes-buch-publiziert/>, (3) Springer Nature, Pressemitteilung, 2.4.2019, <https://www.springer.com/gp/about-springer/media/press-releases/corporateg/springer-nature-maschinen-generiertes-buch/16590072>; (4) Christian Chiarcos im Gespräch mit Monika Seynsche, Deutschlandfunk, 11.4.2019, https://www.deutschlandfunk.de/kuenstliche-intelligenz-erstes-computergeneriertes-buch.676.de.html?dram:article_id=446126; (5) Adrian Lobe: Der Algorithmus als Autor, In: Spektrum.de, <https://www.spektrum.de/kolumne/der-algorithmus-als-autor/1674504>; (6) Madeleine Gregory: AI Trained on Old Scientific Papers Makes Discoveries Humans Missed, Vice, 9.7.2019, https://www.vice.com/en_us/article/neagpb/ai-trained-on-old-scientific-papers-makes-discoveries-humans-missed; (7) David L. Chandler: Can Science writing be automated, MIT News Office, 17.4.2019, <http://news.mit.edu/2019/can-science-writing-be-automated-ai-0418> (6.4.2020).

DIE ROLLE VON KI IN DER WISSENSCHAFT

Zusammenhang mit Megatrends. Nach Ansicht von Prof. Dr. Stefan Wrobel, Leiter des Fraunhofer-Instituts für Intelligente Analyse- und Informationssysteme, sind die Umwälzungen durch KI nur dann richtig zu verstehen, wenn andere Megatrends der Gegenwart mit berücksichtigt werden: „Die Trends der Globalisierung, Digitalisierung und KI hängen zusammen und verstärken sich gegenseitig.“ Die KI sei dabei die entscheidende Ingredienz, da sie die Automatisierung – insbesondere in Bezug auf die Auswertung unstrukturierter Daten – auf eine neue Ebene gehoben habe. Dadurch sei es im Zuge der Globalisierung z.B. überhaupt erst möglich geworden, dass Plattformstrategien immer weitere Bereiche des wirtschaftlichen und gesellschaftlichen Lebens auf der gesamten Welt erfassen und globalisierte Angebote über eine einheitliche Plattform auch lokal erfolgreich sein können.

„Die Trends der Globalisierung, Digitalisierung und KI hängen zusammen und verstärken sich gegenseitig.“

Globale Plattformstrategien. Wrobel berichtete, dass in der Fraunhofer-Allianz Big Data und künstliche Intelligenz mehr als 30 Institute eine branchenübergreifende Expertise bündeln.⁴² Dabei werden aktuelle Entwicklungen in der Wirtschaft deutlich: Auch in Deutschland sind in bestimmten Branchen globale Plattformen schon marktbeherrschend geworden. Dieses Amalgam aus Globalisierung, Digitalisierung und Automatisierung wird nach Wrobel – unabhängig von KI-Hypes – immer weiter wirken. Auch Deutschland müsse sich auf diese Entwicklung einstellen und die Unternehmen in die Lage versetzen, flächendeckend eine KI-getriebene Automatisierung umzusetzen.

In verschiedenen Studien wurde herausgearbeitet, wie viele Arbeitsplätze im Zuge dieses Prozesses möglicherweise vernichtet werden und

⁴² In der Fraunhofer-Allianz Big Data und Künstliche Intelligenz begleiten Fraunhofer-Expert_innen Unternehmen bei der Umsetzung von Big-Data-Strategien, entwickeln Software und datenschutzgerechte Systeme für Big Data und bilden Fach- und Führungskräfte zu „Data Scientists“ aus. Vgl. <https://www.bigdata.fraunhofer.de/de/big-data.html> (6.4.2020).

welche neuen Arbeitsplätze entstehen könnten, wie viel Wertschöpfung wegfallen und neu hinzukommen könnte. An dieser Entwicklung müssten sich Volkswirtschaft und Gesellschaft aktiv beteiligen, meinte Wrobel. Im industriellen Sektor arbeitet Fraunhofer viel mit kleineren und mittleren Unternehmen zusammen, für die der Zugang zu Fachkräften im Bereich KI oft sehr schwierig ist. Deshalb hält Wrobel die Ausbildungs- und Forschungskomponente in der KI-Strategie der Bundesregierung für besonders wichtig. Allerdings müsse man sich bewusst machen, dass KI nur die Speerspitze einer umfassenden, digitalen Umwälzung ist: „Wir erleben eine Revolution des digitalen Gestaltens und des Gestaltens digitaler Prozesse – mit oder ohne KI. KI macht aber vieles erst möglich“, sagte Wrobel.

„Wir erleben eine Revolution des digitalen Gestaltens.“

KI „Made in Deutschland und Europa“. Aufgrund dieser großen Bedeutung von KI findet es Wrobel richtig, dass die Bundesregierung mit ihrer KI-Strategie⁴³ dieses Thema so stark ins Visier genommen hat und mit unterschiedlichen Maßnahmen fördern möchte. Nun sollte darauf geachtet werden, dass die Maßnahmen mit Augenmaß und der nötigen Geduld als langfristige nachhaltige Investition und nicht als Strohfeuer umgesetzt werden. Auch die Zielrichtung sollte immer im Blick behalten werden: „Wo wollen wir technologisch hin, wo genau muss öffentlich investiert werden, damit die gewünschten Entwicklungen in Gang gesetzt werden?“ Erklärtes Ziel sei ja, in Deutschland und Europa eine besondere Form der KI umzusetzen, die ein bestimmtes Gesellschafts- und Wirtschaftsbild reflektiert – im Kontrast zur internationalen Entwicklung, wo stark monopolgetriebene Plattformstrukturen entstehen, die Wirtschaftsbereiche dominieren. Deshalb müsse intensiv über die Frage diskutiert werden, wie die KI-getriebene Ökonomie in Deutschland und Europa als Alternativmodell konkret organisiert werden soll.

Notwendige Maßnahmen. Nach Ansicht von Wrobel sind Investitionen in vielen Bereichen erforderlich, um die gewünschte „digitale Souveränität“ zu erreichen. So müsse Europa in der technologischen Hardware unabhängiger werden, was für die KI besonders wichtig sei. Im Bereich Computer- und Speicherleistungen sollten die begonnenen Investitionen weitergeführt werden, wie z.B. das Big Data Kompetenzzentrum ScaDS Dresden/Leipzig, das vom Bundesministerium für Bildung und Forschung gefördert wird und inzwischen sehr leistungsfähige In-

43 Vgl. https://initiated21.de/app/uploads/2018/01/d21-digital-index_2017_2018.pdf (6.4.2020).

frastrukturen aufgebaut hat.⁴⁴ Auch auf einer übergreifenden Ebene müsse dafür gesorgt werden, dass funktionierende Daten- und KI-Ökosysteme entstehen. „In Deutschland ist die Forschung im Machine Learning und in der KI sehr gut aufgestellt“, sagte Wrobel. Ein Manko sei jedoch, dass die wissenschaftlichen Erkenntnisse den Unternehmen und Anwender_innen nicht flächendeckend zur Verfügung stehen. Deshalb müsse ein übergreifendes Ökosystem entworfen werden, in dem die Ergebnisse aus der deutschen Forschung leicht zertifiziert und für alle verfügbar gemacht werden können.

„In Deutschland ist die Forschung im Machine Learning und in der KI sehr gut aufgestellt.“

Ethische und gemeinwohlorientierte Forschung. PD Dr. Jessica Heesen, Leiterin des Forschungsschwerpunkts Medienethik und Informationstechnik am Internationalen Zentrum für Ethik in den Wissenschaften (IZEW) an der Universität Tübingen, begrüßt das Ziel der Bundesregierung, KI-Forschung und -Anwendung gezielt zu fördern. Die KI-Strategie verfolge richtige Ziele, doch gebe es auch kritische Punkte, die bei einer Umsetzung unbedingt beachtet werden sollten.

Positiv zu bewerten sei, dass in der KI-Strategie die gesellschaftliche und ethische Forschung explizit benannt wird und in Deutschland eine datenschutzkonforme, humanistische und gemeinwohlorientierte KI gestaltet werden soll („KI zum Wohle der Gesellschaft“). Bei genauerer Betrachtung der Maßnahmen zeige sich aber, dass zwar hundert Professuren im KI-Bereich geschaffen werden sollen, aber keine einzige Professur für gesellschaftliche Forschung in Bezug auf KI genannt wird. Ethische und gesellschaftliche Forschung werde nur als dialogische Forschung gefördert, also im Sinne einer Vermittlung und Akzeptanzförderung von KI, nicht jedoch im Bereich der KI-Forschung selbst. Genau an dieser Stelle sei der Bedarf an ethischer und gesellschaftlicher Forschung aber besonders groß. Deshalb sollte bei der Umsetzung der Strategie auch vorgesehen werden, in diesem wichtigen Bereich Forschung zu fördern und Professuren zu schaffen.

„Mit KI werden wesentliche Fragen des Humanen angesprochen.“

Rolle der Geistes- und Sozialwissenschaften. „Mit KI werden wesentliche Fragen des Humanen angesprochen“, sagte Heesen. Deshalb sollten bei der Stärkung und Förderung von KI die Geistes- und Sozialwissenschaften einen hohen Stellenwert erhalten. Die große Bedeutung müsse auch in der konkreten materiellen Ausstattung sichtbar werden. Hier sei noch viel zu tun, auch wenn inzwischen kleine Fortschritte festzustellen sind, z.B. durch die Förderung eines Exzellenzclusters zum Maschinellen Lernen in Tübingen, wo zum Thema Philosophie, Ethik und Maschinelles Lernen geforscht werden wird.⁴⁵ Gegenwärtig sei festzustellen, dass häufig Wissenschaftler_innen über Ethik in der KI sprechen, die keine Ausbildung in diesem Bereich haben und auch keine Forschung dazu betreiben.

Professuren und Nachwuchsförderung. Auf den Plan der Bundesregierung, hundert zusätzliche KI-Professuren zu schaffen, reagierte Heesen mit Skepsis. Der Arbeitsmarkt für hoch qualifizierte Forscher_innen auf diesem Gebiet sei so gut wie leer gefegt – nicht nur auf Ebene der Professuren, sondern auch auf der Ebene wissenschaftlicher Mitarbeiter_innen, die sofort von Unternehmen angeworben werden.

Wrobel meinte, er sei ein großer Befürworter klarer Signale, die dann mit Augenmaß umgesetzt werden sollten. Es sei sicherlich nicht sinnvoll, hundert Professuren innerhalb eines Jahres zu besetzen – und zudem nicht praktikabel, da ausreichend qualifizierte KI-Wissenschaftler_innen in dieser Zahl in dieser kurzen Zeit kaum gefunden werden können. „Ganz entscheidend ist aber das Signal, so etwas tun zu wollen, um in den Institutionen, wo es teilweise große Widerstände gegen fachliche Reorientierungen gibt, etwas in Bewegung zu setzen“, meinte Wrobel. Die große Bedeutung von klaren politischen Signalen sollte nicht unterschätzt werden, da diese deutlich machen, dass ein Aufbruch passieren soll. Schon jetzt sei in vielen deutschen Universitäten dieses Aufbruchssignal angekommen und hätte Veränderungsprozesse ausgelöst. Deshalb wäre es wichtig, an diesem klaren Statement festzuhalten. Gleichzeitig sollte in den Ausführungsbestimmungen dafür gesorgt werden, dass die Universitäten nicht gezwungen sind, diese Positionen kurzfristig innerhalb der nächsten Monate zu besetzen. Dann

45

Auf der Plattform Sci-Hub können wissenschaftliche Aufsätze, die sonst nur hinter einer Paywall online verfügbar sind, illegal heruntergeladen werden. Laut einer Studie von 2017 vermittelt Sci-Hub Zugriff auf 85,2 % aller kostenpflichtigen wissenschaftlicher Aufsätze. Vgl. Achim Wagenknecht: Sci-Hub: Die Guerilla-Bibliothek im Netz, 2.5.2016. In: Unicum.de, <https://www.unicum.de/de/studium-a-z/studium-digital/sci-hub-die-guerilla-bibliothek> (6.4.2020).

wäre es auch möglich, bei der Besetzung der Professuren weiterhin auf Qualität und Exzellenz zu achten und im Zweifel eine Stelle auch erst im nächsten oder übernächsten Jahr zu besetzen.

Nach Wrobel ist es aber unverzichtbar, parallel dazu im Ausbildungsbereich und in der Nachwuchsförderung aktiv zu werden – und zwar in Form einer Pyramide auf allen Ebenen. Es

„Die große Bedeutung von klaren politischen Signalen sollte nicht unterschätzt werden.“

sollten Doktorand_innen- und Postdoc-Programme im Bereich KI gestärkt werden und nicht zuletzt auch der schulische Bereich, indem Schüler_innen für das Programmieren und digitale Gestalten begeistert werden. Da im KI-Programm drei Milliarden Euro investiert werden sollen, seien auch genügend Mittel da, um hundert neue Professuren und gleichzeitig Maßnahmen in der Nachwuchsförderung und im Ausbildungsbereich umzusetzen. Es seien vielfältige Maßnahmen auf unterschiedlichen Ebenen notwendig, die nicht gegeneinander ausgespielt werden sollten.

In der Diskussion im Rahmen der Friedrich-Ebert-Stiftung wurde angemerkt, dass es sinnvoller gewesen wäre, zunächst eine kluge inhaltliche Entwicklungsstrategie für das Land insgesamt zu erarbeiten, daraus dann die Notwendigkeit für bestimmte Forschungsbereiche abzuleiten und diese entsprechend mit Stellen bzw. Professuren zu unterlegen. Erst müsse doch die Frage beantwortet werden, wie und in welchen Bereichen KI inhaltlich entwickelt und gestärkt werden soll und welche Ziele erreicht werden sollen, um entscheiden zu können, in welchen Forschungsbereichen Professuren gebraucht werden.

Aus Wrobels Sicht ist die KI-Strategie – und der darin eingebettete, relativ kleine Bestandteil der hundert neuen KI-Professuren – durchaus das Ergebnis einer intensiv geführten strategischen Debatte. Er berichtete, dass über eine längere Zeit Expert_innenrunden tagten, zahlreiche Papiere verfasst und eine Vielzahl von strategischen Erwägungen berücksichtigt wurden. Ein wichtiger Teil der KI-Strategie seien die Kompetenzzentren zum Thema KI und Maschinelles Lernen – sechs seien schon in Betrieb und weitere sechs sollen noch bundesweit eingerichtet werden. Ein Teil der Professuren könnte dann dort inhaltlich angesiedelt werden. Bei den Forschungsplänen der Kompetenzzentren sei gut zu erkennen, dass ihre Ausrichtung das Ergebnis vieler Debatten war, in welchen Bereichen geforscht werden sollte, z.B. zum Thema erklär-bare KI. Somit beruhe die KI-Strategie durchaus auf einer strategischen Planung und einem übergreifenden inhaltlichen Konzept.

In der weiteren Diskussion wurde vorgeschlagen, die hundert Professuren auf die Länder zu verteilen und die Hochschulen über den Zuschnitt der Professuren entscheiden zu lassen. Dann könnten die Professuren nach dem Profil und den Entwicklungsplänen der Hochschulen in Abstimmung mit den Fakultäten die Stellen mit Wissenschaftler_innen optimal besetzen. Die Erfahrungen mit der Humboldt-Professur würden zeigen, dass die KI-Professuren auch mit qualifizierten internationalen Spitzenwissenschaftler_innen besetzt werden könnten. Mit der von der Humboldt Stiftung finanzierten Humboldt-Professur gelinge es schon seit vielen Jahren, hervorragende Wissenschaftler_innen aller Fächer aus der ganzen Welt für den Wissenschaftsstandort Deutschland zu gewinnen.⁴⁶ Entscheidend

„KI-Forschung muss in der Breite des Wissenschaftssystems stattfinden.“

werde letztlich sein, klare Qualitätskriterien für die KI-Professuren zu benennen und die Stellen konsequent danach zu besetzen.

Spitze und Breite der KI-Forschung. Heesen wies darauf hin, dass die Förderung der KI-Forschung im Zuge der Exzellenzstrategie stark mit einzelnen Leuchttürmen bzw. wenigen Hochschulstandorten verbunden ist. Dies sei jedoch nicht ausreichend. „KI-Forschung muss in der Breite des Wissenschaftssystems stattfinden, da KI die gesamte Gesellschaft betrifft“, sagte Heesen. Universitäten, Hochschulen für Angewandte Wissenschaften und außeruniversitäre Forschungseinrichtungen müssten gleichermaßen daran beteiligt sein, damit die unterschiedlichen Perspektiven auf KI gebührend einbezogen werden können. Häufig finde auch noch eine Verengung auf bestimmte Bereiche der KI-Forschung statt, z.B. Mobilität und Gesundheit. In der Gesellschaft existierten aber noch zahlreiche weitere wichtige Anwendungsbereiche von KI. So spiele KI z.B. in der Bildungsforschung noch keine große Rolle.

46 Inzwischen wurde entschieden, dass bis zu 30 der geplanten KI-Professuren als zusätzliche Alexander von Humboldt-Professuren bis zum Jahr 2024 auf dem Gebiet der Künstlichen Intelligenz besetzt werden sollen. Die Alexander von Humboldt-Professur ist mit 5 Millionen Euro für experimentell und 3,5 Millionen Euro für theoretisch arbeitende Wissenschaftler_innen der höchstdotierte Forschungspreis Deutschlands und wird vom Bundesministerium für Bildung und Forschung finanziert. Jedes Jahr konnten bislang bis zu zehn Humboldt-Professuren verliehen werden. Nun können jährlich sechs weitere Professor_innen speziell für das Gebiet der KI für Deutschland gewonnen werden. Für die neuen Alexander von Humboldt-Professuren für KI können nicht nur Forscher_innen aus technischen Fachgebieten nominiert werden, sondern auch solche, die sich mit den sozio-ökonomischen, ethischen oder rechtlichen Aspekten der KI befassen. Nominierungsberechtigt sind die deutschen Hochschulen. Die Kandidat_innen müssen im Ausland als Forscher_innen etabliert sein. Rund die Hälfte der Preisträger_innen waren bisher deutsche Rückkehrer_innen aus dem Ausland. Vgl. Alexander von Humboldt Stiftung, Pressemitteilung, 8.8.2019, <https://www.humboldt-foundation.de/web/pressemitteilung-2019-15.html> (6.4.2020).

Dominanz eines naturwissenschaftlich-technischen Ansatzes. Auch die kritische Perspektive auf KI müsse gestärkt werden. Wenn sich Deutschland mit einem Alleinstellungsmerkmal im Bereich KI und Ethik und gesellschaftliche Forschung profilieren möchte, müssten hier deutlich mehr Maßnahmen in der Breite umgesetzt und mehr finanzielle Mittel investiert werden. In der Exzellenzstrategie zur Förderung von Spitzenforschung würden 57 Exzellenzcluster ab 2019 für mindestens sieben Jahre gefördert – und nur fünf davon gehörten zu den Geistes- und Sozialwissenschaften, einschließlich der Digital Humanities. „Bisher ist der Zugang zu KI-Fragestellungen sehr stark natur- und technikwissenschaftlich geprägt. Die kritische Reflexion von KI bleibt oft auf der Strecke“, sagte Heesen.

„Die kritische Reflexion von KI bleibt oft auf der Strecke.“

Qualitative Unterschiede bei KI-Forschung. Typischerweise würden zumeist die positiven Aspekte und Potenziale von KI benannt, etwa Sprachassistenzsysteme für Blinde oder neue Bilddiagnoseverfahren bei der Brustkrebserkennung. Bei diesen Themen sei in der Regel kein Widerspruch zu erwarten bzw. die Reaktionen seien durchweg positiv, weil damit neue Möglichkeiten für die Menschen verbunden sind. Der Übergang zu problematischen Formen der KI verschwimme dann häufig. So sei z.B. von Apple ein KI-System vorgestellt worden, bei dem Menschen alle Daten über ihre alltäglichen Interaktionen sammeln und auch per Kamera aufzeichnen können. Ein solches System werde als Möglichkeit beworben, die eigene Gedächtnisleistung zu verbessern: Alles, was ein Mensch macht, kann erfasst werden, und Apple verfüge über die Software zur Entschlüsselung der Daten. So wird es z.B. möglich, die Aufzeichnung eines Gesprächs Monate später abzurufen. Mit solchen Systemen würden umfassendere Fragen tangiert, auch der mögliche Missbrauch von Daten. „Das ist der Umschlagpunkt: Hier brauchen wir ethisch-geisteswissenschaftliche Forschung, die feststellt, dass das nicht mehr das Gleiche ist wie ein Sprachassistenzsystem“, sagte Heesen. Bei solchen Verfahren kämen fundamentale qualitative Unterschiede zum Tragen. Es sei ein anderer Zugriff notwendig, der nicht über KI-Forschung selbst geleistet werden könne, sondern nur über eine Metaperspektive und kritische Reflexion.

Deshalb wäre es aus Sicht von Heesen auch wichtig, die geplanten zusätzlichen hundert Professor_innenstellen für KI fachlich breit anzulegen. Mit der Einrichtung der KI-Professuren werde zwar ein klares politisches Signal gesendet, doch stelle sich die Frage, welches Wissenschaftsver-

ständnis damit gestärkt werde. Von KI-Programmen profitiere meistens die Forschung in den Technik- und Naturwissenschaften – und damit ein Wissenschaftsverständnis, das stark von Quantifizierung und der Herstellung von Korrelationen geprägt ist. Der reflexive Ansatz der Geistes- und Sozialwissenschaften rücke dabei in den Hintergrund, obwohl klar sei, dass mit einem naturwissenschaftlich-technischen Verständnis manche Entwicklungen nicht analysiert oder erklärt werden können.

In der Diskussion wurde zugestimmt, dass die Geistes- und Sozialwissenschaften in Verbindung mit KI-Forschung einen hohen Stellenwert haben sollten, doch dürften sie in Kooperationen nicht zum „Reparaturbetrieb“ oder zur „Bremse“ werden. Vielmehr sollten sich Geistes- und Sozialwissenschaftler_innen überlegen, was sie aktiv zum Thema beitragen können. So würden in Diskussionen über KI z.B. viele Begriffe aus der europäischen Kultur- und Geistesgeschichte relativ unreflektiert verwendet, etwa Intelligenz, Lernen, Verstehen und Erklären. Diese Begriffe seien in der Philosophie bereits sehr differenziert analysiert und definiert worden. Die Geisteswissenschaften könnten hier für mehr Klarheit sorgen, aber auch in anderen Bereichen wichtige Beiträge leisten. Wissenschaftler_innen dieser Fachrichtungen sollten eine offensivere Rolle als bisher einnehmen. Dabei sollte möglichst auf den Begriff der Begleitforschung verzichtet werden, weil dadurch der Eindruck erweckt werde, die Analyse würde nebenbei stattfinden, ohne Einfluss auf den Forschungsprozess nehmen zu können.

Datenökonomie. Eine weitere wichtige Frage betrifft die Datenbasis für KI. Die EU strebt einen freien Datenbinnenmarkt an. Welche Voraussetzungen sind dafür erforderlich? Daten sind die unverzichtbare Grundlage für das Training von KI-Systemen, da diese mit Wissen versorgt werden müssen. Eine entscheidende Frage ist dabei nach Wrobel, wie zukünftig die Verknüpfung und Verfügbarkeit von Daten wirtschaftlich organisiert werden soll. Als marktwirtschaftlich organisierte Volkswirtschaft mit sozialer Verantwortung habe die Bundesrepublik stark von einer kleinteiligen und wettbewerblich organisierten Wirtschaft profitiert. Im Bereich der Daten würden derzeit aber auf globaler Ebene völlig andere Strukturen entstehen, indem Daten monopolisiert und auf große Plattformen gelegt werden, die dann große Bereiche der Weltwirtschaft dominieren. Ziel sollte es jedoch sein, dass unterschiedliche Akteure der Gesellschaft und der Wirtschaft nach ihren Vorstellungen und Werten Daten miteinander austauschen können.

Die EU will Datenökonomie auch mit Datensouveränität verbinden – nicht nur für Unternehmen, sondern auch für einzelne Verbraucher_innen. Fraunhofer arbeitet daran, die aktuellen technischen Möglichkeiten der Informatik so zu nutzen, dass neue Arten der Datenökonomie entstehen können. Der Ansatz des Data Space (Datenraumarchitektur) zeichnet sich durch zahlreiche Teilnehmer_innen aus: Vielfältige Unternehmen und Einzelpersonen stellen Daten bereit und/oder nutzen sie. Die Daten werden kryptografisch gesichert, quasi „in Umschläge gepackt“, und nur zertifizierte Anwender_innen können dann auf diese Daten für bestimmte Zwecke zugreifen. Auf diese Weise könnte nach Wrobel eine Ökonomie entstehen, die nach den Prinzipien einer Marktwirtschaft organisiert ist – mit einer Vielzahl an Anbieter_innen und Nutzer_innen von Daten. Dies sei die Grundidee des Data Space-Konzeptes.

„Der Ansatz des Data Space zeichnet sich durch zahlreiche Teilnehmer_innen aus.“

An der International Data-Space-Association (ISDA) sind derzeit schon mehr als hundert Unternehmen und Forschungseinrichtungen beteiligt. Die ISDA hat schon Mitglieder aus 17 Ländern und in fünf Ländern sogar offizielle Repräsentanzen (Partnerorganisationen, die diese Rolle in ihrem Land übernehmen). Wichtig wäre aus Wrobels Sicht, schon jetzt dafür zu sorgen, dass möglichst viele verschiedene Akteure in solch ein ökonomisches System KI-Apps und KI-Technologien einbringen können, sodass in den wettbewerblich betriebenen App-Stores KI-Leistungen föderal und marktwirtschaftlich gemeinsam erbracht werden. Damit könnte dem gegenwärtig dominierenden System entgegengewirkt werden, in dem einzelne große Plattformen – ob staatlich oder privatwirtschaftlich – die Daten besitzen und alle anderen nur Kunden dieser Plattform sind. Dies ist nach Wrobel die Grundfrage: Schaffen wir es, ein System der Datenökonomie zu organisieren, in dem es eine Vielzahl von Partner_innen gibt, die miteinander flexibel und nach eigenen Vorstellungen kooperieren können und sowohl Daten bereitstellen als auch Daten nutzen? Die Zertifizierung spielt dabei eine zentrale Rolle.

Zertifizierung im Anwendungskontext. Fraunhofer entwickelt gerade Zertifizierungskriterien für KI. Einzelne Kriterien können noch nicht benannt werden, da der Prozess noch nicht abgeschlossen ist. Es seien aber viele Unternehmen und Organisationen bereit, sich zertifizieren lassen, berichtete Wrobel. Man müsse sich immer bewusst machen, dass nicht am KI-System selbst erkannt werden kann, ob es gut oder richtig, schlecht oder falsch ist. Am Grundalgorithmus eines Tiefen Neuralen Netzes könne z.B. nicht abgelesen werden, ob die KI vertretbar

handele. Dies könne nur sichergestellt werden, wenn die Prozesse einer Organisation oder eines Unternehmens, die KI betreiben, zertifiziert werden. Bei KI sei – wie bei anderen technologischen Entwicklungen auch – Missbrauch nicht völlig auszuschließen. Deshalb müssten KI-Systeme immer im Anwendungskontext betrachtet und zertifiziert werden. Entscheidend seien dabei

- die Daten, die verwendet werden,
- die Vorkehrungen, die getroffen werden, um sicherzustellen, dass die Daten richtig sind,
- die Vorkehrungen für Reaktionszeiten, um Fehlentwicklungen zu unterbinden etc.

Auf dieser Basis werde auch ein großer Unterschied erkennbar zwischen Organisationen und Ländern, die KI verantwortungsvoll, erklärbar, nachvollziehbar, transparent und gerecht einsetzen und Organisationen und Ländern, die das nicht tun.

POLITIK UND KI – EINE GEREGETE BEZIEHUNG?

Breiter Ansatz der KI-Strategie. René Rösper, Bundestagsabgeordneter und Sprecher der SPD-Bundestagsfraktion in der Enquete-Kommission „Künstliche Intelligenz – Gesellschaftliche Verantwortung und wirtschaftliche Potenziale“ des Deutschen Bundestags, machte deutlich, dass die hundert Professuren als Botschaft verstanden werden sollten und die Länder durch Bundesmittel entlastet werden. Die konkrete Ausgestaltung könne erst nach den Gesprächen mit den Ländern festgelegt werden. Offen sei auch noch die Frage, wie diese Stellen auf Dauer finanziert werden.

Rösper betonte, dass die KI-Strategie relativ breit und auch flexibel aufgestellt ist. Dazu gehöre z.B. die Einbindung der Alexander von Humboldt-Stiftung zur Gewinnung qualifizierter internationaler Expert_innen für die KI-Professuren, die Stärkung der Nachwuchsausbildung in den Kompetenzzentren, aber auch die Ausschreibung eines offenen Professurenprogramms für die Hochschulen. Rösper findet es wichtig, dass bei der Besetzung der hundert Professuren auf Diversität geachtet wird, also auch weibliche und junge Wissenschaftler_innen ausreichend zum Zug kommen. Die Professuren sollten dazu genutzt werden, hier gezielt Impulse zu setzen. Die Besetzung der KI-Professuren könnte durchaus sukzessive stattfinden und daran angepasst werden, ob genügend qualifizierte Wissenschaftler_innen vorhanden sind.

Stärkung der Geistes- und Sozialwissenschaften. Die SPD-Fraktion teile die Auffassung von Heesen, dass die Geistes- und Sozialwissenschaften im KI-Bereich gestärkt werden müssen. Die Förderung von KI-Forschung dürfe sich nicht auf Wissenschaftler_innen aus Mathematik, Physik, Ingenieurwesen etc. beschränken. Er habe die Botschaft verstanden, dass auch KI-Professuren gebraucht werden, die sich dezidiert mit ethischen Fragen in der KI-Forschung beschäftigen und kritische Reflexion ermöglichen. Wichtig sei dabei auch die Gemeinwohlorientierung der KI, die in der KI-Strategie festgehalten wurde.

Europäisches Konzept. Die großen Unterschiede zur KI in China und den USA sind nach Rösper nicht im Bereich der Technologie zu finden, die im

*„Welche Gesellschaft wollen wir
und welchen Beitrag kann KI
dazu leisten?“*

Prinzip überall gleich ist. Entscheidend sei vielmehr die Zielsetzung und die zentrale Frage: „Welche Gesellschaft wollen wir und welchen Beitrag kann KI dazu leisten?“ Das in Europa verfolgte Konzept

von KI, das auf ein demokratisches und vielfältiges Gemeinwesen ausgerichtet ist, werde letztlich länger tragen als die Konzepte in den USA oder in China, wo KI-Forschungsprogramme sich an sicherheitspolitischen bzw. militärischen Fragen orientieren oder in ein autoritäres Gesellschaftskonzept eingebettet sind.

Venture Capital. Waltinger hatte in seinem Vortrag angesprochen, dass es in Deutschland zwar viele kreative und engagierte junge KI-Wissenschaftler_innen gibt, viele aber nicht im Land gehalten werden können, weil hier die notwendige Venture-Capital-Szene als Risikokapitalgeber nicht vorhanden ist, wie z.B. in den USA. Nach Ansicht von Waltinger ist Venture Capital in ein ganzes Ökosystem eingebettet, bei dem es darum geht, junge und innovative Unternehmen ausreichend und angemessen zu fördern – auch unabhängig von KI. Deutschland und Europa könnten in der Forschung zwar große Erfolge im Bereich Maschinelles Lernen vorweisen, doch sei das in der Öffentlichkeit kaum bekannt. Das liege unter anderem daran, dass die großen internationalen Leuchtturmprojekte im Bereich KI von großen Unternehmen gefördert werden, die Talente aus den Universitäten abziehen und ihnen ein attraktives Arbeitsumfeld sowie umfangreiche finanzielle Mittel bieten, die Universitäten ihnen nicht geben können. Das Ungleichgewicht zwischen Unternehmen und Universitäten sei in diesem Bereich beträchtlich. Deshalb sollte die Politik überlegen, wie neben guter industrieller Forschung auch die universitäre Forschung gestärkt und jungen Talenten die notwendige Infrastruktur und Ausstattung gegeben werden kann.

Röspel stellte heraus, dass die öffentliche Hand in Bezug auf finanzielle Mittel nicht mit den großen Playern der Wirtschaft mithalten kann. Aus seiner Sicht müsse die Vorgehensweise in Deutschland anders sein als z.B. in den USA. Deshalb werden in Deutschland auch andere Maßnahmen ergriffen, z.B. wurden die Mittel des Programms EXIST – Existenzgründungen aus der Wissenschaft⁴⁷ mehr als verdoppelt und es soll eine Agentur für Sprunginnovationen eingerichtet werden. Sicherlich sei auch

47 Vgl. <https://www.exist.de/DE/Home/inhalt.html> (6.4.2020).

eine andere Haltung erforderlich: Es brauche mehr Mut und mehr Akzeptanz für das Scheitern, das bei innovativen Wegen keine Seltenheit sei. In diesem Bereich ist in Deutschland nach Ansicht von Röspel zwar ein langsamer Veränderungsprozess festzustellen, doch sei dies nicht mit der Situation in den USA vergleichbar, weil hier eine andere Mentalität vorherrscht. Dennoch sollte darüber nachgedacht werden, wie das Umfeld für Innovationen weiter verbessert und talentierte Menschen darin bestärkt werden können, etwas zu wagen.

Verantwortung bei KI-Systemen. Eine wichtige Frage bei KI-Systemen ist, wer für die Prozesse die Verantwortung trägt, insbesondere bei Fehlern. Prof. Dr. Thomas Wischmeyer, Professor der Rechtswissenschaften an der Universität Bielefeld und Mitglied der Datenethikkommission⁴⁸ vertritt die Auffassung, dass aus rechtlicher Sicht die erforderlichen Instrumente im Prinzip vorhanden sind, mit deren Hilfe Verantwortung zugewiesen werden kann. Der Bundestag könnte etwa auch für KI eine Gefährdungshaftung beschließen und festlegen, dass „X“ – z.B. das Unternehmen, der Plattformbetreiber, der Nutzer – die Verantwortung im rechtlichen Sinne trägt. Ob der dahinterstehende Verantwortungsbeitrag dann immer adäquat ist, auch mit Blick auf die Innovationsfähigkeit Deutschlands, sei eine andere Frage, die letztlich politisch entschieden werden müsse.

„Im ganzen Bereich Software und Internet findet eine praktische Verantwortungsdiffusion statt.“

Bei der aktuellen Debatte über die Verantwortung von KI-Systemen ist nach Wischmeyer die interessante Frage, wie Verantwortung aufgefasst wird: Verantwortung im Recht werde in der Regel typischerweise sehr kleinteilig geregelt. Nun sei eine Veränderung zu erkennen: „Im ganzen Bereich Software und Internet findet eine praktische Verantwortungsdiffusion statt“, sagte Wischmeyer. Wenn Systeme Fehler machen, würden sich daraus verschiedene Probleme ergeben:

1. Da die Systeme enorm komplex und sehr unterschiedliche Akteure daran beteiligt sind, stehe der/die Einzelne vor der Frage, wer in Haftung genommen werden kann. Diese Frage sei häufig sehr schwer zu beantworten.
2. Der/die Einzelne müsse den Nachweis erbringen, dass ein Fehler stattge-

48

Vgl. <https://www.bmi.bund.de/DE/themen/it-und-digitalpolitik/datenethikkommission/datenethikkommission-node.html> (6.4.2020).

funden hat. Hier stelle sich die Frage, woher er/sie das wissen kann, insbesondere dann, wenn die Systeme sehr komplex und schwer erklärbar sind.

3. Der/die Einzelne müsse erst einmal wissen, dass das System einen Fehler gemacht hat.

„Die Herausforderung für den Gesetzgeber besteht darin, für Haftungsklarheit zu sorgen – sowohl im allgemeinen Bereich der Software als auch im speziellen Bereich KI und autonome Fahrzeuge“, sagte Wischmeyer. Wer letztlich den entscheidenden Verantwortungsbeitrag geleistet hat, sei aus Sicht der Anwender_innen weniger interessant als zu wissen, wer praktisch zur Verantwortung gezogen werden kann. Um für Haftungsklarheit zu sorgen,

„Die Herausforderung für den Gesetzgeber besteht darin, für Haftungsklarheit zu sorgen.“

müsse deshalb der Gesetzgeber eine Person bzw. eine konkrete Stelle festlegen, an die sich der oder die Einzelne wenden kann.

Smart Contracts. Diskutiert wurde auch über die Frage, ob für Deutschland Gefahren durch Smart Contracts⁴⁹ für das Rechtssystem ausgehen könnten, insbesondere durch die amerikanische Variante, die stark technikgetrieben auf dem Vormarsch ist. In den USA kann es z.B. passieren, dass eine Person, die mit einer Mietzahlung im Rückstand ist, keinen Zugang mehr zu ihrer Wohnung bekommt, weil diese automatisch verriegelt wird. Hier stellt sich die Frage, ob ein solcher Mechanismus dem deutschen Rechtssystem widersprechen würde, weil der Fall einseitig definiert wird und der Beklagte kein Anhörungsrecht hat. Wischmeyer wies darauf hin, dass es zwar einzelne Anwendungsfälle, aber noch keinen flächendeckenden Einsatz von Smart Contracts gibt. Die dahinterstehende, grundlegende Frage sei dabei, wie weit das Recht „vollautomatisiert“ werden soll. Bisher hätten staatliche Gerichte bzw. Menschen die Klärung in Streitfragen übernommen und dabei technologische Möglichkeiten genutzt, um sich auf das Wesentliche konzentrieren zu können. In Deutschland soll die Streitentscheidung einem Menschen aus unterschiedlichen Gründen überlassen werden – nicht jedoch, weil ein Mensch objektiver entscheidet. Schließlich sei aus der Forschung bekannt, dass Richter_innen häufig sehr subjektiv urteilen. Der wesentliche Grund sei vielmehr,

49 Ein Smart Contract ist ein „elektronischer Vertrag, der hinterlegte Regeln automatisch überwacht und definierte Aktionen bei Vorliegen eines Trigger-Events selbstständig ausführen kann.“ (Andreas Mitschele: Definition Smart Contract, Gabler Wirtschaftslexikon, <https://wirtschaftslexikon.gabler.de/definition/smart-contract-54213> (6.4.2020)).

dass ein gesellschaftliches Bedürfnis nach einer menschlichen Entscheidung bei der Streitentscheidung besteht.

Hinter dem Trend der Smart Contracts steht nach Wischmeyer ein entdifferenziertes Gesellschaftskonzept: Wenn früher eine Person ihre Miete nicht bezahlte, ging der Fall ins Rechtssystem und wurde dort bearbeitet. Bei Smart Contracts – am krassensten beim chinesischen Sozialkreditsystem – würden Sanktionen extrem schnell in der Lebenswelt spürbar. Damit höre das Rechtssystem auf, ein selbstständiges konfliktmediatierendes System zu sein, wo nach einer längeren Zeit und rechtlichem Gehör eine Entscheidung getroffen wird, die lebensweltliche Konsequenzen hat. Bei Smart Contracts würden die Menschen sofort von bestimmten Aspekten ausgeschlossen bzw. spürten unmittelbare Konsequenzen. „Diese gesellschaftsentdifferenzierende Tendenz muss genau beobachtet werden“, sagte Wischmeyer. Solche Lösungen könnten zwar den Eindruck von Effizienz hervorrufen, doch seien sie in Deutschland mit den Strukturrentscheidungen, auf die das gesellschaftliche System aufgebaut ist, kaum kompatibel. Wichtig ist nach Wischmeyer, bei Smart Contracts zu unterscheiden, ob Verbraucher_innen davon betroffen sind oder ob sie etwa im Bereich des Finanzderivathandels angewendet werden. Im Business-to-Business-Bereich könnten Smart Contracts durchaus einen Sinn haben und seien in der Regel problemlos.

Heesen merkte dazu an, dass in der neuen EU-Datenschutzgrundverordnung ein Passus vorhanden ist, nach dem Menschen niemals Objekt einer alleine automatisierten Entscheidung sein dürfen. Aus dieser Perspektive wäre somit ausgeschlossen, dass eine Person automatisch keinen Zugang mehr zu ihrer Wohnung hat. Demnach sei es das Recht jedes Menschen, dass immer ein Mensch beurteilt, ob eine Entscheidung legitim ist.

Autonom fahrende Autos. In öffentlichen Diskussionen werden immer wieder autonom fahrende Autos thematisiert, insbesondere in Bezug auf die Frage, wie ein solches KI-System in Konfliktsituationen mit Menschen auf der Straße agieren würde. Wenn der Algorithmus eine Differenzierungsmöglichkeit und Entscheidungshierarchie enthalten würde, könnte das zu ethisch problematischen Entscheidungen des KI-Systems führen. Wie sollte mit dieser Gefahr umgegangen werden und wer würde in solchen Fällen die Verantwortung tragen?

Das Thema Verantwortung bei autonom fahrenden Autos kann nach Ansicht von Wischmeyer an sich mit dem bisherigen Haftungsrecht gelöst werden, weil dieses technikneutral angelegt ist. Das BGB stamme aus dem Jahr 1900 und habe bisher noch alle technologischen Wandlungen über-

standen. Allerdings sei es durchaus hilfreich, bei solchen grundlegenden Innovationen wie autonom fahrenden Autos neue gesetzliche Regelungen zu schaffen, in denen ein konkreter Anspruch formuliert und speziell bestimmt wird, wer die Verantwortung tragen soll. Das würde den Bürger_innen mehr rechtliche Klarheit geben. Darüber hinaus sei die Frage der Verantwortung beim autonomen Fahren durch die Expertenkommission unter Di Fabio behandelt worden.⁵⁰ Diese Kommission habe klar herausgearbeitet, dass zwischen dem Wert einzelner Menschen nicht differenziert werden dürfe.

Nach Auffassung von Schmid wird an dieser Frage sehr deutlich, dass Ethik nicht im Kompetenzbereich der KI-Forscher_innen, sondern bei Ethiker_innen und Philosoph_innen liegt, die schon sehr lange zur Entscheidungsfindung in moralischen Dilemmatasituationen forschen. Auch Wischmeyer betonte, dass technologische Entwicklungen immer durch kritisch reflektierende und ethische Forschung ergänzt werden müssen. Die Datenethikkommission habe auch die Funktion, eine „ethische Folie“ in politische Prozesse einzubringen, was bei der KI-Strategie auch geschehen sei. Ihre Ausrichtung auf ethische und gemeinwohlorientierte KI sei maßgeblich auf eine Stellungnahme der Datenethikkommission zurückzuführen.

Bei komplexen Wertschöpfungsketten in der Industrie 4.0 geht es nach Wischmeyer auch um die wichtige Frage, wie der dabei generierte Wohlstand verteilt werden soll. Damit seien politische Fragen eingeschlossen, die eine rechtliche und ethische Dimension haben. In der Datenethikkommission stehe man vor der Herausforderung, dass es bisher relativ wenig belastbare sozialwissenschaftliche und philosophische Forschung zu diesen Grundlagenfragen gibt. Deshalb arbeite die Kommission daran, in diesem Bereich Leitlinien für die Bundesregierung zu formulieren. Dabei habe man festgestellt, dass es kein Problem sei, Sachverständige zu gewinnen, die Auskunft über die neuesten Technologien oder den Transfer in die Industrie geben können. Man könne aber nur auf einen sehr kleinen Pool von Geistes- und Sozialwissenschaftler_innen zurückgreifen, die über diese Fragen forschen und sich einbringen können. Deshalb halte er die Forderung von Heesen für sehr wichtig, die Geistes- und Sozialwissenschaften bei der Umsetzung der KI-Strategie und insbesondere der Ausrichtung der KI-Professuren stärker zu berücksichtigen.

⁵⁰ Vgl. Bundesministerium für Verkehr und digitale Infrastruktur (Hrsg.): Automatisiertes und Vernetztes Fahren. Bericht der Ethik-Kommission, Juni 2017, https://www.bmvi.de/SharedDocs/DE/Publikationen/DG/bericht-der-ethik-kommission.pdf?__blob=publicationFile (6.4.2020).

Ethische Fragen und Werte. Nach Ansicht von Wischmeyer sind im Bereich von KI noch viele hoch relevante gesellschaftliche Fragen zu bearbeiten. Grundsätzlich sei mit dem rechtlichen Rahmen in Deutschland ein bestimmtes ethisches Konzept verbunden. Die in der Verfassung festgelegten Werte wie Demokratie und Menschenwürde müssten auch bei jeder neuen Technologie beachtet werden. Teilweise würden durch neue Technologien aber auch neue Fragen entstehen, bei denen viele noch keine ethische Intuition entwickelt haben. Dazu gehöre z.B. auch die Frage, wie weit personalisierte Profilbildung bei Daten zugelassen werden soll.

Röspel sieht hier das Kernproblem: Wenn man versuche, Verhaltensregeln für Maschinen aus menschlicher Betrachtung zu formulieren, stelle man fest, dass die Menschen oft gar keine Klarheit über Verhaltensregeln und ihre ethischen Grundlagen besäßen. Ethik und Geisteswissenschaften seien unverzichtbar, um solche grundlegenden Fragen aufzuarbeiten. Dilemmata-Konfliktsituationen seien schon von Menschen kaum zu entscheiden bzw. lösbar, z.B. bei der Organtransplantation. Trotzdem werde oft erwartet, Maschinen so zu programmieren, dass sie sich eindeutig zur Zufriedenheit der Menschen entscheiden können. Es ginge um Grundsatzfragen, die gesellschaftlich geklärt werden müssen, wenngleich dies in einer Demokratie nicht final gelingen könne und solle. „Am Ende müssen juristische und politische Entscheidungen getroffen werden“, sagte Röspel.

„Am Ende müssen juristische und politische Entscheidungen getroffen werden.“

Umgang mit Daten und Datenschutz. Wischmeyer merkte an, dass sich in juristischen Debatten der Fokus vom Algorithmus hin zu den Daten verschiebt, weil dort die grundlegenden Entscheidungen zu treffen sind und ethische und rechtliche Fragen von eminenter Bedeutung aufgeworfen werden. Entscheidend sei dabei die Frage der Trennung zwischen personenbezogenen und nicht personenbezogenen Daten: Allerdings sei es durchaus schwierig, Daten im rechtlichen Sinne zu anonymisieren. Vor dem Hintergrund der bisherigen Rechtsprechung des Europäischen Gerichtshofes müsse zudem festgestellt werden, dass zahlreiche Daten, die im Unternehmens- oder Forschungskontext anfallen, als personenbezogene Daten zu werten sind. Entsprechend seien in beiden Bereichen die Anforderungen aus dem Datenschutzrecht zu beachten, auch wenn die Freiräume für die Forschung deutlich größer sind als für Unternehmen.

In Bezug auf die Daten muss nach Wischmeyer stark differenziert werden: Welche Art von Daten sind im System? Wie weit soll die Offenlegungspflicht gehen? Sollen die Rohdaten oder die aggregierten Daten betroffen

sein etc.? Dabei können noch verschiedene Stufen unterschieden werden: Wem gegenüber soll die Offenlegungspflicht gelten? Soll z.B. ein zertifiziertes Unternehmen mit einem Schlüssel Zugriff darauf haben? Oder sollen bestimmte Daten allgemein zugänglich sein (Open Data)? Auch müsse man sich Gedanken darüber machen, welche Datenarten für die jeweilige Nutzung und die Art des Zugangs in Frage kommen. Ergänzend komme dann noch die wichtige Souveränitätsfrage hinzu: Wie viel wollen wir in Deutschland öffentlich machen, was auch außerhalb Europas genutzt werden kann? Enteignet sich Deutschland als Forschungsstandort selbst, wenn die Zugangsrechte sehr weitgehend sind?

Zur Plattformökonomie merkte Wischmeyer an, dass sich aus dem Zusammenschluss großer Unternehmen auf einer Plattform möglicherweise Konstellationen ergeben, die nicht mit dem Kartell- und Wettbewerbsrecht vereinbar sind. An dieser Frage arbeite gegenwärtig eine weitere Kommission „Wettbewerb 4.0“ beim Bundeswirtschaftsministerium, teilweise auch die Datenethikkommission. „Wir brauchen ein Update für unser Wettbewerbsrecht, um mit solchen großen Plattformen umgehen zu können.“ Darüber hinaus sei damit auch eine übergreifende Gerechtigkeitsfrage verbunden: Wie weit wollen wir auf unsere Daten Zugriff haben, wo sehen wir volkswirtschaftliche Potenziale, wo Gemeinwohlpotenziale?

„Wie gehen wir mit Daten um und wer verfügt über sie?“

Röspel verwies auf das SPD-Positionspapier „Daten-für-alle-Gesetz“⁵¹, das Datenmonopole verhindern will und die Zielsetzung verfolgt, Daten als Gemeingut zu sehen. Es müsse darüber nachgedacht werden, wem die Daten gehören sollen: demjenigen, der sie generiert oder demjenigen, über den sie erhoben werden? Der bekannte Spruch „Meine Daten gehören mir“ sei zwar einfach und nachvollziehbar, doch sei das Problem komplexer und es müssten ein wichtige Punkte geklärt werden: Wie gelingt es dem Staat einerseits, das Individuum vor dem Missbrauch seiner Daten zu schützen, und andererseits zu erreichen, Daten zum Gemeinwohl zu nutzen? Hier sei der Reflexionsprozess noch nicht abgeschlossen, da in Daten große Potenziale stecken, ihre Nutzung aber auch mit großen Gefahren verbunden ist. Die zentrale Frage laute: „Wie gehen wir mit Daten um und wer verfügt über sie?“ Transparenz, Offenheit und Gemeinwohlorientierung seien hierbei unverzichtbar.

51 Vgl. <https://www.spd.de/aktuelles/daten-fuer-alle-gesetz/> (6.4.2020).

- #13 Angela Borgwardt: **Digitalisierung in der Wissenschaft** (2018)
- #12 Angela Borgwardt: **Impulse für die strategische Debatte in der Wissenschaft** (2017)
- #11 Angela Borgwardt: **Neuer Artikel 91B GG – Was ändert sich für die Wissenschaft** (2015)
- #10 Angela Borgwardt: **Wissenschaftsregionen – Regional verankert, global sichtbar** (2015)
- #09 Angela Borgwardt: **Wissenschaft auf Abwegen? Zum drohenden Qualitätsverlust in der Wissenschaft** (2014)
- #08 Angela Borgwardt: **Leitlinien des zukünftigen Wissenschaftssystems – Grundforderungen, Gemeinsamkeiten und Widersprüche** (2014)
- #07 Angela Borgwardt: **Europäische Forschungsallianzen – Regionale Verbünde und EU-Förderung** (2013)
- #06 Angela Borgwardt: **Internationaler, besser, anders? – Die Strukturen des Wissenschaftssystems nach 2017** (2012)
- #05 Angela Borgwardt: **Internationalisierung der Hochschulen – Strategien und Perspektiven** (2012)
- #04 Angela Borgwardt: **Rankings im Wissenschaftssystem – Zwischen Wunsch und Wirklichkeit** (2011)
- #03 Angela Borgwardt: **Der lange Weg zur Professur – Berufliche Perspektiven für Nachwuchswissenschaftler/innen** (2011)
- #02 Angela Borgwardt, Marei John-Ohnesorg: **Vielfalt oder Fokussierung – Wohin steuert das Hochschulsystem nach drei Runden Exzellenz?** (2010)
- #01 Meike Rehbarg: **Verbündete im Wettbewerb – Neue Formen der Kooperation im Zuge der Exzellenzinitiative, dargestellt am Beispiel des Karlsruher Instituts für Technologie** (2007)

Das **Netzwerk Wissenschaft** entwickelt Beiträge und Empfehlungen zur künftigen Gestaltung des deutschen Wissenschaftssystems.

Die Publikationen können Sie per E-mail nachbestellen bei: marion.stichler@fes.de

Digitale Versionen aller Publikationen:

<http://www.fes.de/themen/bildungspolitik/index.php>



Besuchen Sie unseren Bildungsblog
www.fes.de/bildungsblog

ISBN: 978-3-96250-563-9

1. Auflage

Copyright by Friedrich-Ebert-Stiftung

Hiroshimastraße 17, 10785 Berlin

Abt. Studienförderung

Redaktion: Dr. Martin Pfafferoth, Marion Stichler, Lena Bülow

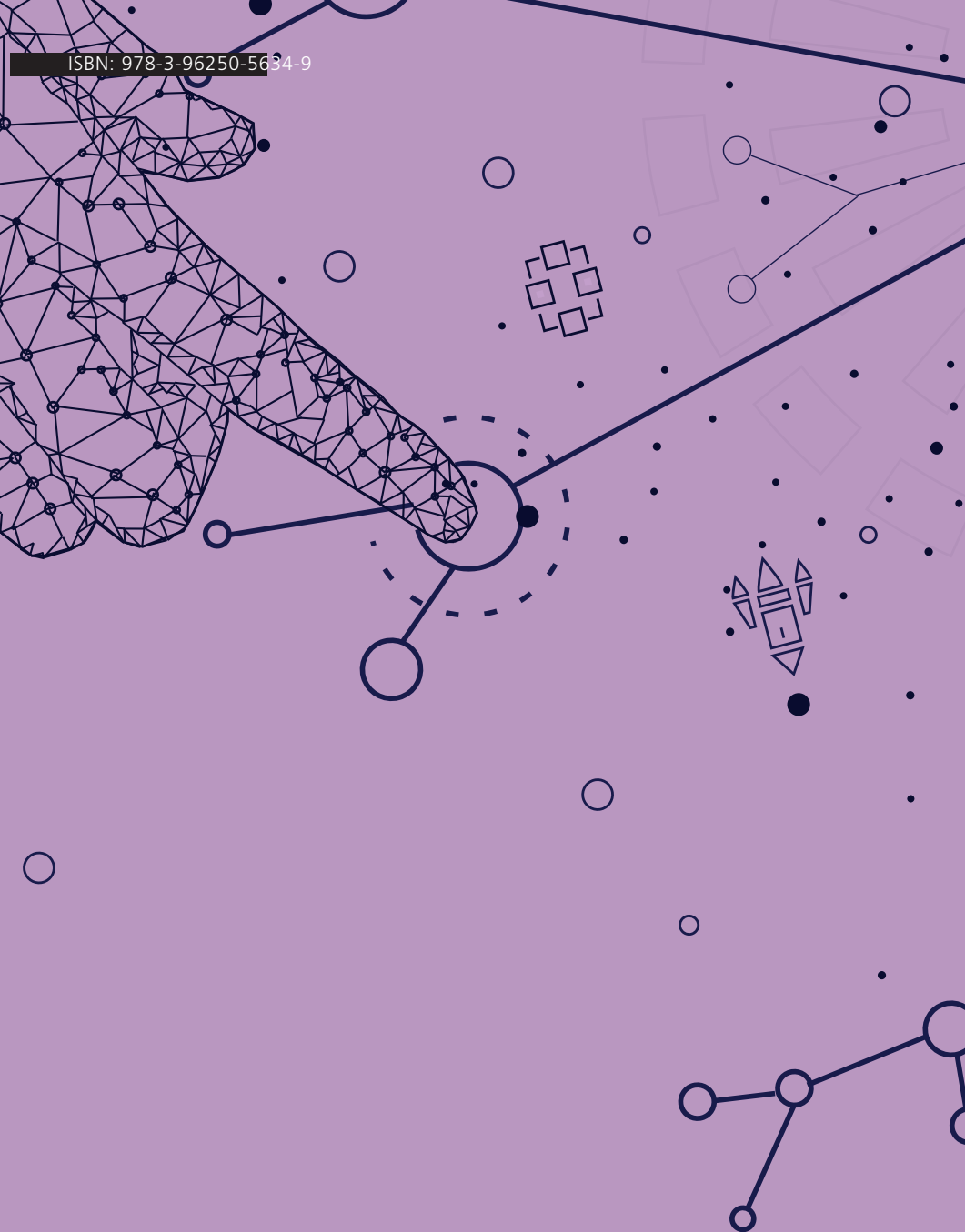
Satz & Umschlaggestaltung: minus Design, Berlin

Coverbild: © Vladimir Vihrev/Shutterstock

Druck: Druckerei Brandt GmbH, Bonn

Printed in Germany 2020

ISBN: 978-3-96250-563-9



Committed to excellence

Die Friedrich-Ebert-Stiftung ist im Qualitätsmanagement zertifiziert nach EFQM
(European Foundation for Quality Management): Committed to Excellence