

KÜNSTLICHE INTELLIGENZ & ETHIK

Robotergeretze

Schon in den 1940er-Jahren formulierte Science-Fiction-Autor Isaac Asimov (Bild) eine Reihe von „Robotergeretzen“, damit Menschen die Kontrolle über die Künstliche Intelligenz behalten können.

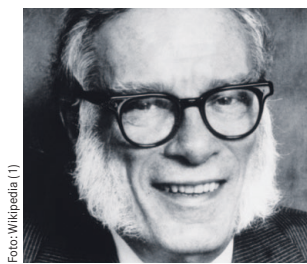


Foto: Wikipedia (1)

Algorithmen

Ethik-Experten pochen auf das „Recht auf Erklärung“: Programmierer sollen offenlegen, auf Basis welcher Kriterien der Algorithmus eines Computersystems Entscheidungen trifft.

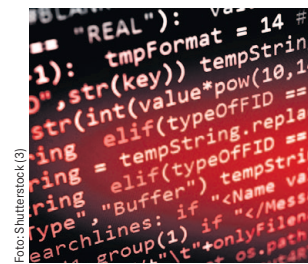


Foto: Shutterstock (3)

Seinen evolutionären Siegeszug verdankt der Mensch seiner herausragenden Intelligenz. Heute wird die Evolution durch Künstliche Intelligenz vorangetrieben – aber wohin führt das?

Der nächste SPRUNG



Von Sonja Bettel

Künstliche Intelligenz. Der Name sagt schon, worum es geht, nämlich darum, die menschliche Intelligenz nachzubilden. Konkret bedeutet das, Technologien zu entwickeln, die eigenständig Probleme bearbeiten können. Während das Ziel einer intelligenten Maschine bisher nur in Science Fiction-Filmen existiert, sind Teilbereiche der Künstlichen Intelligenz, abgekürzt KI, bereits Realität.

Stefan Woltran, Professor am „Institute of Logic and Computation“ der Technischen Universität Wien, spricht lieber von KI-Technologie als von „Künstlicher Intelligenz“. Denn der Begriff sei irreführend, wie er am Beispiel eines Bonmots erklärt: „Wir KI-Forscher sagen, die Humanbiologie habe so viel mit KI zu tun wie die Ornithologie mit der Luftfahrttechnik.“ Zwar könnte man sagen, dass Flugzeuge nach dem Vorbild des Vogels entwickelt worden seien, aber sie haben keine Federn und schlagen nicht mit den Flügeln. Andererseits können sie sehr schnell fliegen, was ein Vogel nicht kann. „Nach diesem Bild ist das, was wir machen, Düsenantriebsforschung“, sagt Woltran schmunzelnd. Er und seine Kolleginnen und Kollegen beschäftigen sich nämlich damit, Algorithmen für bestimmte Aufgaben besser und schneller zu machen.

Sprachverarbeitung im Aufwind

Obwohl KI-Technologie ein weites Feld ist und noch viel Forschung bedarf, ist sie in unserem Alltag bereits omnipräsent. Zum Beispiel in jedem Smartphone, erläutert Woltran, der selbst keines besitzt, in jeder Google-Suchmaschine, jeder maßgeschneiderten Werbeanzeige auf YouTube oder Kaufempfehlung auf Amazon. Besonders spannend sei in den letzten Jahren die Entwicklung der Sprachverarbeitung für das Übersetzen oder Generieren von Texten, oder auch Sprachassistenten wie Alexa, Siri, Cortana oder Bixby.

Es sei beeindruckend, was diese Systeme bereits leisten, anerkennt der KI-Forscher, aber man dürfe nicht glauben, dass sie Sprache verstehen. Sie haben einfach nur aus einer riesigen Menge an Texten und dem Regelwerk der Grammatik gelernt, was sie aufgrund einer konkreten Frage tun beziehungsweise was sie antworten sollen. Den Unterschied zum menschlichen Sprachverständnis erklärt Stefan Woltran so: „Die Frage ‚Kann ein Krokodil an einem Hürdenlauf teilnehmen?‘ kann eine KI vermutlich nicht beantworten. Der Mensch hingegen kann sich ein Krokodil mit seinen kurzen

Beinen und Hürden bildlich vorstellen und so rasch die Antwort finden: Nein, kann es nicht.“ Ein anderes Problem wurde mittlerweile technisch gelöst, zumindest zum Teil: Schon ein Kind kennt den Unterschied zwischen einem Hund und einer Katze, für Software war das schwierig. Nach dem Training mit Millionen von annotierten Bildern („das ist ein Hund, das ist eine Katze ...“) hat das Programm den Unterschied gelernt, auch wenn man ihm ein bisher unbekanntes Bild zeigt. Hier sieht man aber auch die Probleme der KI: Bilderkennungsprogramme hätten auch schon bei einem völlig abstrakten Bild behauptet, das sei zu 99 Prozent ein Hund, sagt Woltran.

Wie eine KI-Technologie zu ihren Ergebnissen kommt, weiß man bei derartigen lernenden Systemen oft nicht. Wenn die Kreditwürdigkeit eines Bankkunden, die Gewaltbereitschaft eines Besuchers im Fußballstadion, die von einem Reisenden ausgehende Terror-

„Rufe nach einem bedingungslosen Grundeinkommen, einer ‚Robotersteuer‘ oder einer grundlegend anderen Schulbildung kommen sogar von jenen, die mit Technologie ein Vermögen verdienen.“

Gefahren durch KI

Erneut warnt ein Expertenbericht vor dem Missbrauch Künstlicher Intelligenz durch skrupellose Regierungen, Kriminelle und Terroristen („The Malicious Use of Artificial Intelligence“, Feb. 2018; <https://maliciousaireport.com>).

gefahr und vieles mehr quasi in einer Blackbox errechnet wird und der Mensch nicht weiß, wie das Ergebnis zustande kommt und inwieweit es stimmt, stehen wir jedoch vor großen Problemen. Das hat auch die Europäische Union erkannt und deshalb in der EU-Datenschutz-Grundverordnung, die am 25. Mai 2018 in Kraft tritt, vorgesehen, dass Bürgerinnen und Bürger ein Recht auf Information haben. Wenn Firmen nicht erklären können, wie ihre lernenden Maschinen die Kundendaten verarbeiten, drohen hohe Strafen. Die KI-Forschung ist nun gefordert, diese Blackbox transparent zu machen – was nicht einfach wird.

Politisch einflussreiche „Chatbots“

Wichtig wäre das auch für „Soziale Netzwerke“. In den letzten Monaten wurde immer klarer, dass in ihnen „Chatbots“, also nicht-menschliche Nutzer, aktiv sind, die die politische Meinungsbildung wesentlich beeinflussen können. Für die menschlichen Nutzer ist es oft sehr schwer, zwischen Mensch und Ma-

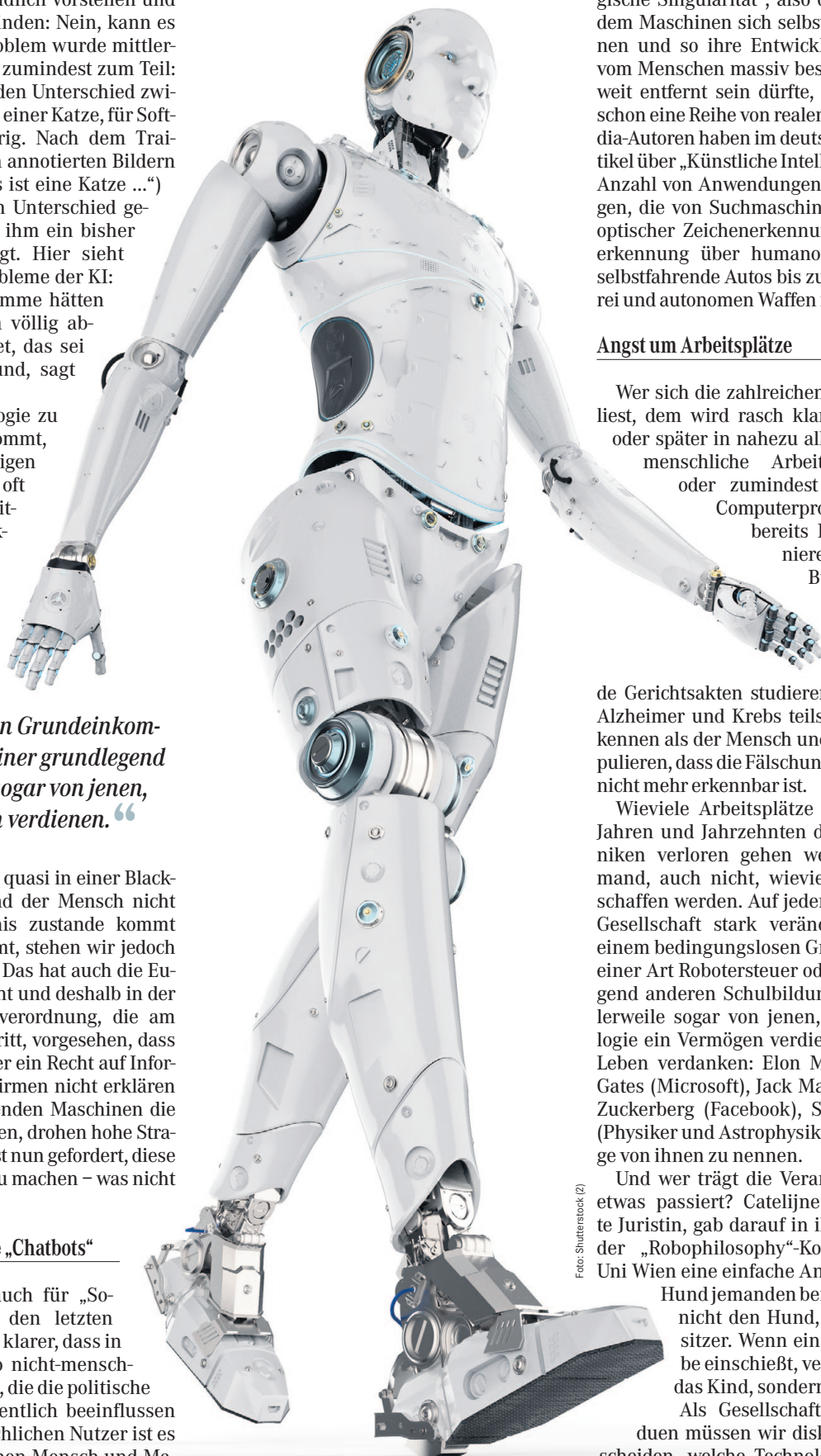


Foto: Shutterstock (2)

schine unterscheiden zu können, weshalb wiederholt eine Kennzeichnungspflicht für KI gefordert wird. Das sei so ähnlich wie mit dem Straßenverkehr Anfang des 20. Jahrhunderts, bemerkt Woltran. Bei der Einführung von Autos und Straßenbahnen sei es zu Chaos und schweren Unfällen gekommen. So wurde klar, dass es Ampeln und Verkehrsregeln braucht. Catelijne Muller, Mitglied des Europäischen Wirtschafts- und Sozialausschusses und Rapporteur des aktuellen EU-Meinungsbildungsprozesses über Künstliche Intelligenz und Gesellschaft, fordert, dass die Menschen nicht die Herrschaft über die Technik aufgeben dürfen. Sie propagiert dafür sogar einen eigenen Hashtag: #humanincommand.

Auch wenn die viel diskutierte „Technologische Singularität“, also der Zeitpunkt, ab dem Maschinen sich selbst verbessern können und so ihre Entwicklung unabhängig vom Menschen massiv beschleunigen, noch weit entfernt sein dürfte, birgt die KI jetzt schon eine Reihe von realen Risiken. Wikipedia-Autoren haben im deutschsprachigen Artikel über „Künstliche Intelligenz“ eine große Anzahl von Anwendungen zusammengetragen, die von Suchmaschinen, Übersetzung, optischer Zeichenerkennung und Gesichtserkennung über humanoide Roboter und selbstfahrende Autos bis zu Medizin, Juristerei und autonomen Waffen reicht.

Angst um Arbeitsplätze

Wer sich die zahlreichen Beispiele durchliest, dem wird rasch klar, dass KI früher oder später in nahezu allen Bereichen die menschliche Arbeitskraft ersetzen oder zumindest ergänzen wird. Computerprogramme können bereits Popsongs komponieren, Bilder malen, Bücher schreiben, Unfallschäden an Autos begutachten, in kurzer Zeit Tausen-

de Gerichtsakten studieren, Unkraut jäten, Alzheimer und Krebs teils verlässlicher erkennen als der Mensch und Videos so manipulieren, dass die Fälschung mit freiem Auge nicht mehr erkennbar ist.

Wieviele Arbeitsplätze in den nächsten Jahren und Jahrzehnten durch diese Techniken verloren gehen werden, weiß niemand, auch nicht, wieviele durch sie geschaffen werden. Auf jeden Fall wird KI die Gesellschaft stark verändern. Rufe nach einem bedingungslosen Grundeinkommen, einer Art Robotersteuer oder einer grundlegend anderen Schulbildung kommen mittlerweile sogar von jenen, die mit Technologie ein Vermögen verdienen oder ihr das Leben verdanken: Elon Musk (Tesla), Bill Gates (Microsoft), Jack Ma (Alibaba), Mark Zuckerberg (Facebook), Stephan Hawking (Physiker und Astrophysiker) – um nur einige von ihnen zu nennen.

Und wer trägt die Verantwortung, wenn etwas passiert? Catelijne Muller, studierte Juristin, gab darauf in ihrem Vortrag bei der „Robophilosophy“-Konferenz an der Uni Wien eine einfache Antwort: „Wenn ein Hund jemanden beißt, verklagen wir nicht den Hund, sondern den Besitzer. Wenn ein Kind eine Scheibe einschleift, verklagen wir nicht das Kind, sondern die Eltern.“

Als Gesellschaft und als Individuen müssen wir diskutieren und entscheiden, welche Technologien wir wollen

Autonome Autos

Fahrerlose Autos sind bereits in Testung. Das zwingt zu Entscheidungen bei der Programmierung, zum Beispiel: Wer wird bei einem Unfall geschützt, die Insassen oder die Passanten auf der Straße?



„Killerdrohnen“

Wissenschaftler aus aller Welt engagieren sich für ein Verbot autonomer Waffen: Militärische Technologien könnten in falsche Hände geraten und zur unberechenbaren Bedrohung werden.



„Mangelndes Vertrauen gibt es immer, wenn Menschen mit Systemen interagieren, die sie letztlich nicht durchschauen können.“

und welche nicht, mahnt Mark Coeckelbergh, Technikphilosoph an der Universität Wien. Wichtig ist diese Diskussion auch in einem Bereich, der eher abseits der Öffentlichkeit stattfindet: jenem der Waffen. Kampfroboter im Stile eines Terminator aus dem gleichnamigen Film hält er für Science-Fiction, tatsächlich seien Soldaten aber heute schon mit viel High-Tech ausgestattet und mit Drohnen können Militärs bereits über Tausende Kilometer Entfernung gezielt auf Menschen schießen. Um deutlich zu machen, wohin die Entwicklung gehen könnte, haben kritische KI-Wissenschaftler aus aller Welt ihre Forderung nach dem Verbot autonomer Waffen vor etwa drei Monaten mit einem drastischen

„Computerprogramme können bereits Popsongs komponieren, Bilder malen, Bücher schreiben, Unfallschäden begutachten, Diagnosen stellen oder in kurzer Zeit Tausende Gerichtsakten studieren.“

Video illustriert. Darin dringen „Slaughterbots“ – Drohnen, kleiner als eine Handfläche – in eine Universität ein und töten gezielt politisch aktive Studenten. Heute ist das noch Fiktion, aber wie lange noch?

Warum wir autonom agierende Waffen verbieten sollen, erklärt Philosoph Coeckelbergh so: „Wir Menschen erleben Angst vor Schmerz und Tod und wissen, dass andere Menschen dies ebenfalls empfinden. Das führt dazu, dass wir normalerweise zögern, wenn wir einen anderen verletzen oder töten könnten. Die Maschine hat dieses Empfinden nicht.“ Aus diesem Grund sollten wir sicherstellen, dass der Mensch die letzte Entscheidungsgewalt niemals einer „Künstlichen Intelligenz“ überlässt. Derzeit tun wir das aber Schritt für Schritt durch die Hintertür.

Wie sehr sind wir bereit, uns der Künstlichen Intelligenz anzuvertrauen? Viele Menschen könnten ihrem Einsatz mit Angst und Skepsis begegnen.

Der gefühlte Kontrollverlust

Von Martin Tauss

Stellen Sie sich vor, Sie sind Arzt oder Ärztin in einer Krebsklinik. Um Ihre Arbeit zu unterstützen, wird Ihnen von der Klinikleitung ein Super-Computer zur Seite gestellt, der als „Wunderwuzzi“ der Krebstherapie angepriesen wird – nennen wir ihn „Watson“, wie das berühmte Computerprogramm von IBM. Für die zwölf häufigsten Krebsarten liefert er auf Knopfdruck Behandlungsempfehlungen nach den neuesten wissenschaftlichen Erkenntnissen. Doch ganz wohl fühlen Sie sich nicht mit dem obergescheiten Watson. Wenn er für eine Therapie plädiert, die mit Ihrer Einschätzung übereinstimmt, verspüren Sie ein Gefühl größerer Sicherheit, sehen aber keinen großen Wert in Watson: Denn die Künstliche Intelligenz (KI) sagt Ihnen lediglich, was Sie ohnehin schon wissen.

Der Arzt aus der Box

Schlägt Watson hingegen eine andere Therapie vor, als Sie zu verschreiben gedenken, zweifeln Sie eher an seiner als an Ihrer Kompetenz. Wer weiß denn schon, wie die Maschine zu ihrem Schluss gekommen ist: Die Algorithmen sind zu komplex, als dass sie für Menschen völlig nachvollziehbar wä-



„Mehr Transparenz, wie die Algorithmen eines selbstlernenden Systems funktionieren, würde die Einstellung zur Künstlichen Intelligenz verbessern.“

ren. Als Spezialist vertrauen Sie lieber auf Ihr handfestes Wissen und Ihre langjährige Erfahrung.

„Watson for Oncology“ ist angetreten, um das Überleben von Krebspatienten zu verbessern. Das System kommt in mehr als 150 Krankenhäusern in elf Ländern zum Einsatz. Doch die Nutzung funktioniert nicht reibungslos; die erwähnten Probleme gibt es im klinischen Alltag tatsäch-

lich. Zwei Krankenhäuser haben unlängst die Kooperation mit dem KI-System beendet, weil viele Ärzte schlicht nicht auf Watson vertrauten. Das berichtet der Computerforscher Vyacheslav Polonski vom Oxford Internet Institute, der sich darüber Gedanken macht, wie das Vertrauen in die Künstliche Intelligenz gestärkt werden kann. Denn ähnliche Probleme gibt es immer, wenn Menschen mit Systemen interagieren, die sie nicht durchschauen können.

Vertrauensfördernde Maßnahmen

Wer schon Erfahrungen mit Künstlicher Intelligenz gemacht hat, bringt diesen Systemen größeres Vertrauen entgegen, so Polonski. Auch mehr Transparenz, wie die Algorithmen funktionieren, würde die Einstellung verbessern. Und es wäre gut, die Menschen nicht ganz aus dem maschinellen Entscheidungsprozess auszuschließen: Gibt es die Möglichkeit, den Algorithmus leicht zu modifizieren, steigt die Zufriedenheit mit der Nutzung von Künstlicher Intelligenz, betont der Forscher.

Freilich empfiehlt sich auch ein Mindestmaß an kritischem Denken: Das ist eine Eigenschaft, die man auch sonst gut brauchen kann. Denn im Gegensatz zum IQ trägt es nachweislich zur allgemeinen Lebenszufriedenheit bei.

ARBEITEN FÜR WIEN

WIR GEBEN KINDERN
HALT, WENN IHRE WELT
INS WANKEN GERÄT.

Entgeltliche Einschaltung

Manchmal passiert etwas in einer Familie und Kinder müssen ganz schnell weg von zuhause. Als **Krisenpflegeeltern** haben Sie Erfahrung mit Kindern, sind zeitlich flexibel und schenken Kindern aller Altersstufen liebevolle Aufnahme. Die Stadt Wien bietet Ihnen als Krisenpflegeeltern mit Hauptwohnsitz in Wien neue Anstellungsmodelle und Weiterbildungen.

Wenn Sie sich dieser Herausforderung stellen wollen und bereit zur Zusammenarbeit sind, informieren Sie sich beim MAG ELF Servicetelefon 01 4000 8011 und auf www.pflegemama.at bzw. www.pflegepapa.at

**WIR SUCHEN
KRISENPFLERGEELTERN**

Stadt Wien