

ULRICH
EBERL

**33 FRAGEN →
ANTWORTEN**

KÜNSTLICHE INTELLIGENZ

PIPER

**33 FRAGEN →
ANTWORTEN**

Künstliche Intelligenz ist jetzt schon allgegenwärtig – im Smartphone und darüber hinaus: Gesichtserkennung, Sprachassistenten, personalisierte Werbung, Übersetzungsprogramme, Gesundheits-Apps. Schnell wird sich diese Technologie tief in unser Leben eingraben: Wir werden Fahrzeuge auf Autopilot schalten, Hand in Hand mit Robotern arbeiten und immer mehr lästige Aufgaben an Maschinen übertragen. Doch auch die Gefahren sind unübersehbar: der perfekte Überwachungsstaat, verdeckte Manipulationen, autonome Kampfmaschinen, drohende Jobverluste. Können Maschinen vielleicht sogar intelligenter werden als Menschen? Ist Künstliche Intelligenz eher ein Horrorszenario oder doch die größte Verheißung der Menschheit? Antworten auf die wichtigsten Fragen bietet dieses Buch.

ISBN 978-3-492-31578-4



9 783492 315784

www.piper.de

€ 10,00 [D]
€ 10,30 [A]

PIPER

Inhalt

Einleitung	7
1. Was ist Künstliche Intelligenz, und was unterscheidet sie von menschlicher Intelligenz?	9
2. Warum fasziniert uns Künstliche Intelligenz so sehr?	13
3. Warum gibt es gerade jetzt einen solchen KI-Boom?	15
4. Ist Künstliche Intelligenz die größte technische Revolution seit 250 Jahren?	18
5. Funktioniert Künstliche Intelligenz wie unser Gehirn?	21
6. Was versteht man unter Deep Learning?	26
7. Auf welchen Feldern ist Künstliche Intelligenz bereits besser als jeder Mensch?	28
8. Kann es Intelligenz ohne einen Körper geben?	34
9. Was können Roboter mit Künstlicher Intelligenz heute leisten?	36
10. Wie menschenähnlich können Roboter werden?	39
11. Kann man Roboter wie kleine Kinder erziehen?	42
12. Wird Künstliche Intelligenz unseren Alltag revolutionieren?	45
13. Wo liegen die Einsatzgebiete bei Unternehmen, Umwelttechnik und Infrastrukturen?	49
14. Sind Maschinen die besseren Ärzte?	53
15. Sind USA und China in Künstlicher Intelligenz uneinholbar vorn?	58

16. Ist Künstliche Intelligenz nur etwas für reiche Staaten?	63
17. Wird ein Großteil aller Jobs verloren gehen?	66
18. Was sollen Kinder heute lernen, und welche neuen Berufe wird es morgen geben?	70
19. Maschinen machen unsere Arbeit – ist das nicht wie im Schlaraffenland?	73
20. Werden wir in einer Gemeinschaft von Menschen und smarten Maschinen leben?	76
21. Bedeutet Künstliche Intelligenz das Ende der Privatsphäre?	79
22. Kann uns Künstliche Intelligenz perfekt manipulieren?	85
23. Wie gefährlich werden die Hacker von morgen?	88
24. Wird der Krieg der Zukunft von Robotern entschieden?	91
25. Kann Künstliche Intelligenz Fehler machen und Vorurteile haben?	94
26. Wie bekommt man Moral in die Maschine?	97
27. Kann eine Maschine Emotionen erkennen oder gar selbst Gefühle haben?	101
28. Können Maschinen kreativ sein?	105
29. Kann es Maschinen mit Bewusstsein geben?	108
30. Kann eine Superintelligenz die Menschheit vernichten?	111
31. Werden Cyborgs entstehen – Mischungen aus Mensch und intelligenter Maschine?	114
32. Wie wird sich die Künstliche Intelligenz bis 2050 weiterentwickeln? ...	117
33. Wo liegen die Grenzen der Künstlichen Intelligenz?	121
Literatur	125

Was ist Künstliche Intelligenz, und was unterscheidet sie von menschlicher Intelligenz?



Bei Künstlicher Intelligenz (KI) geht es darum, »Maschinen dazu zu bringen, Dinge zu tun, die – würden sie von Menschen vollbracht – Intelligenz erfordern würden«. So hat es vor über 50 Jahren einer der Pioniere auf diesem Feld, Marvin Minsky, formuliert, und bis heute gibt es keine treffendere Definition. Denn letztlich führt uns dies auf die Frage zurück, was Intelligenz an sich ist. Im Allgemeinen gilt sie als die Fähigkeit, Probleme zu lösen und mit neuen Anforderungen zurechtzukommen. In diesem Sinne sind auch Tiere bis zu einem gewissen Grad intelligent: etwa Schimpansen, die Steine als Hammer und Amboss nutzen, um Nüsse zu knacken, oder Kraken, die Gläser mit Schraubverschluss öffnen, oder Delfine, die im Team neue Jagdtechniken entwickeln.

Allerdings würde die Messung eines Intelligenzquotienten bei Tieren grandios scheitern. Denn IQ-Tests konzentrieren sich auf die Vervollständigung von Zahlenreihen, räumliches Vorstellungsvermögen, Bilderrätsel, Textaufgaben, Sprachverständnis und logisches Denken. Was bei Tieren nicht funktioniert, das klappt nun aber bei Maschinen immer besser: Sie lernen Objekte erkennen, Sprache verstehen und Texte lesen. Sie können Wissen verarbeiten und Schlussfolgerungen ziehen. Sie gewinnen beim Spiel der Könige, beim Schach, und beim Brettspiel Go gegen menschliche Weltmeister, sie bluffen beim Pokern, spüren Krebszellen auf und optimieren die Routen von Lieferfahrzeugen in Echtzeit – bald könnten smarte Maschinen in IQ-Tests besser abschneiden als Menschen.

Doch macht dies Maschinen klüger als uns? Keineswegs, denn Intelligenz ist nicht nur das, was der IQ misst. Zum Lösen vieler realer Probleme braucht man mehr als kognitive Fähigkeiten: neben der räumlichen, sprachlichen, mathematischen und logischen Intelligenz zum Beispiel auch die emotionale und soziale Intelligenz, wenn man mit anderen Individuen kooperieren muss. Nicht zu vergessen die sensorische Intelligenz koordinierter Bewegungen. Denn genau genommen, sagen Experten, sei es schwieriger, beim Schachspiel die Figuren zu greifen und zu setzen, als das Spiel zu gewinnen. Und in der Tat: Der menschliche Schachweltmeister wurde schon 1997 von einem Computer entthront. Wissenschaftlern ist es aber erst unlängst – nach vielen Jahren der Forschung – gelungen, Robotern beizubringen, wie man eine Tür mit einem Schlüssel öffnet. Wir halten Schach und Go für intellektuell herausfordernd, weil man dafür Abstraktionsvermögen und strategisches Denken braucht, doch für Maschinen ist das eher eine leichte Übung.

Der alte Spruch »Computern und Robotern fällt leicht, was Menschen schwerfällt – und umgekehrt« galt jahrzehntelang als unumstößliche Wahrheit. Computer können blitzschnell große Zahlen multiplizieren, und Maschinen können tonnenschwere Lasten heben, aber Türen öffnen und Bälle fangen, sicher laufen und Hindernissen ausweichen, das ist für sie extrem schwierig. Für gesunde Menschen ist all das kein Problem, aber auch wir haben lange gebraucht, um solche Fähigkeiten zu erwerben – wir nehmen sie nur als selbstverständlich wahr, weil vieles unbewusst abläuft. Babys lernen zunächst krabbeln, laufen und greifen, erst nach Jahren kommen Sprechvermögen, Vorstellungskraft und die Fähigkeit, Handlungen zu planen, hinzu. Eine Tür gezielt mit einem Schlüssel zu öffnen – das schaffen Kinder erst im Alter von etwa vier Jahren.

Um eine solche Leistung richtig bewerten zu können, muss man sich nur anschauen, wie viele Nervenzellen mit der Feinabstimmung der Motorik beim Gehen, Sprechen und Greifen beschäftigt sind. Im Kleinhirn, wo dies stattfindet, befinden sich dicht gepackt mehr als zwei Drittel aller Nervenzellen eines erwachsenen Menschen – deutlich mehr als in der Großhirnrinde, die für die bewussten Denkvorgänge zuständig ist. Wenn ein Roboter nach einer Teetasse greifen soll, dann wird schnell klar, wie viel sensomotorische Intelligenz dafür nötig ist: die Koordination von Auge und Hand, das Feedback der Tastsensoren, das sofortige Nachsteuern, um die Tasse nicht kippen zu lassen, und vieles mehr.

Neben Sehen, Sprechen, Lesen und der Verarbeitung großer Datenmengen sollten leistungsfähige KI-Maschinen also auch sensomotorische Intelligenz besitzen. Auf all diesen Gebieten haben Forscher in den letzten Jahren erhebliche Fortschritte erzielt – weit mehr als in den Jahrzehnten zuvor. Wir befinden uns daher tatsächlich am Beginn des Zeitalters der smarten Maschinen. Bereits heute verwenden wir viele KI-Verfahren ganz selbstverständlich, etwa im Smartphone: Entsperren per Gesichtserkennung, Sprachassistenten wie Siri, Alexa & Co., automatische Textkorrekturen, Bildersuche im Internet, der News Feed in sozialen Medien, personalisierte Werbung, Übersetzungsprogramme, Gesundheits-Apps – all das funktioniert mit KI. Auf den Straßen fahren erste vollautomatische Fahrzeuge, in Fabriken arbeiten Roboter Hand in Hand mit Menschen, in Kliniken und Banken beraten Computer Ärzte und Finanzfachleute, und täglich kommen neue Anwendungen hinzu.

Dennoch: Dies sind erst Maschinen mit einer sogenannten schwachen Intelligenz. Sie sind sehr gut auf speziellen Feldern, aber ohne Allgemeinintelligenz. Sie sind lernfähig und werden immer besser darin, konkrete Anwendungspro-

bleme zu lösen. Von einer starken KI, die der unseres Gehirns gleicht oder es sogar übertrifft – und die vielleicht ein eigenes Bewusstsein entwickeln könnte –, sind heutige Maschinen allerdings noch weit entfernt.

Denn dafür müssten sie in der Lage sein, aus ihren Sensor-signalen Modelle von ihrer Umgebung und sich selbst zu erstellen, sie müssten neugierig sein, ihre Umwelt erforschen, Wünsche verspüren, Gefühle empfinden, sich selbst Ziele setzen, mit anderen in Beziehung treten, kommunizieren – kurz: zu einer eigenständigen, moralisch verantwortlichen Persönlichkeit werden.

Nichts deutet darauf hin, dass dies Maschinen in absehbarer Zeit gelingen könnte. Doch das ist auch gar nicht nötig. Schon die Fortschritte der Künstlichen Intelligenz, die sich heute abzeichnen, sind so vielfältig und umfassend, dass sie unsere Welt radikal verändern werden – in allen Bereichen unseres Lebens.

Je erfolgreicher KI-Systeme in Zukunft in Büros und Banken, Kliniken und Fabriken, auf der Straße und zu Hause funktionieren, desto größer wird die Gefahr, dass Menschen zu sehr auf Maschinen vertrauen und ihre Empfehlungen ungeprüft übernehmen. Doch das wäre fatal, wie etliche Beispiele zeigen: Vor einigen Jahren hat ein Bilderkennungssystem von Google Fotos dunkelhäutiger Menschen in die Kategorie »Affen« einsortiert. Eine Software von Amazon, die die besten Jobbewerber finden sollte, hat systematisch Frauen benachteiligt, und eine KI der US-Justizbehörden hat für Afroamerikaner doppelt so häufig wie für Weiße fälschlich eine hohe Rückfallquote bei Straftaten vorhergesagt.

Was war hier falsch gelaufen? Die Bild-Software hatte offenbar bei ihren Lernbeispielen die Farbe höher gewichtet als andere Kriterien, die Menschen von Affen unterscheiden. Die Probleme der KI-Systeme für Personal und Justiz hingegen sind subtiler, weil sie Vorurteile, die bereits in den Trainingsdaten enthalten sind, implizit übernommen haben. Amazon hatte in der Vergangenheit weit mehr Männer eingestellt als Frauen, und so kam der Algorithmus zum Schluss, dass Bewerbungen von Frauen grundsätzlich schlechter zu bewerten seien. Ähnliches galt für die US-Behörden: Dort wurden die Bewertungskriterien oft von Weißen erstellt.

Daher sollte jeder Anwender die Grenzen von KI-Systemen kennen. Insbesondere funktionieren die am meisten verwendeten Deep-Learning-Netze rein nach Verfahren der mathe-

mathischen Optimierung (siehe Frage 6). Dafür brauchen sie viele Tausend Trainingsdaten, doch wie sie zu ihren Ergebnissen kommen, ist kaum transparent. Eine Überprüfung durch Menschen ist daher unerlässlich. Außerdem lassen sich solche Netze gezielt täuschen, etwa durch raffiniert hinterlegte Strukturen. So entdeckten Bildverarbeitungssysteme in Fotos, die für Menschen nur Rauschen enthielten, Gürteltiere oder Pfauen. Japanische Wissenschaftler konnten sogar zeigen, dass im Extremfall der Austausch eines einzigen Bild-Pixels genügt, damit ein Deep-Learning-Netz ein Pferd plötzlich mit 99,9 Prozent Sicherheit als Frosch tituliert.

In einem Projekt namens DeepDream präsentierten Forscher einem Rechner, der mit Tierbildern trainiert worden war, nur Fotos von Wolken. Was passierte? Der Algorithmus entdeckte darin unter anderem Fische, Hunde und Vögel – wie ein Kind, das in Wolken Fabelwesen hineininterpretiert. Als die Forscher die Rückkopplung im Netz verstärkten, wurden die Muster, die der Algorithmus fand und sichtbar machte, immer verrückter: Am Himmel erschienen farbige Strudel wie bei van Gogh, und aus Gehsteigen quollen bunte Kühe hervor. Das Ganze wirkte wie LSD-Bilder – eine Ähnlichkeit, die nicht zufällig ist. Denn auch im Gehirn kommt es bei Halluzinationen zu einer verstärkenden Rückkopplung, einer Überinterpretation von Signalen. Unter Drogeneinfluss sieht und hört man Dinge, die nicht da sind – genau dies passiert auch Deep-Learning-Netzen, die unbedingt Muster finden sollen.

Was diesen Netzen offensichtlich fehlt, ist das Hintergrundwissen – etwa dass es in Wolken keine Fische gibt. Wer derartige KI-Systeme nutzt, muss sich also bewusst machen, dass sie nicht über Alltagsintelligenz verfügen, sondern nur Muster herausfiltern und mathematische Zusammenhänge herstellen. Das Problem dabei: Beim Finden von Mustern blendet man gleichzeitig die Nicht-Muster, die Abweichungen, aus. Eine

solche KI hat sozusagen eine eingebaute Tendenz, Minderheiten kleinzuhalten. Wenn sie Fehler macht und diese nicht erkannt werden, besteht die Gefahr, dass sich die KI im Lauf der Zeit stets selbst bestätigt und noch fehleranfälliger wird.

Natürlich gelten auch für die Entwickler von KI Recht und Gesetz: Antidiskriminierung, Produkthaftung, Gewährleistung, Strafrecht – hier darf es keine Abstriche geben. Es darf nicht passieren, dass ein Roboter im Hotel Menschen je nach Hautfarbe, Geschlecht oder Alter unterschiedlich behandelt. Oder dass dies gar ein selbstfahrendes Auto in einer Unfallsituation tut. Nehmen wir eine Bank als Beispiel: Ein Programm, das die Kreditwürdigkeit von Personen bewerten soll, sortiert vielleicht Menschen mit bestimmten Vornamen aus, weil es eine Korrelation zwischen Vornamen und der Bedienung von Krediten gefunden hat. Wenn nun aber ein Kunde den Verdacht hat, dass er aus einem derartigen Grund benachteiligt wurde, hat er ein Recht auf Auskunft. Wenn die Maschine keine vernünftige Erklärung liefern kann, warum der Kredit abgelehnt wurde, muss ein Mensch eingeschaltet werden, der die Maschine überstimmen darf – und in diesem klaren Fall von Diskriminierung auch überstimmen muss.

→ 26.

Wie bekommt man
Moral in die Maschine?

Entscheidend für lernende Maschinen ist – wie bei Menschen auch – die Qualität der Lehrer. Ein misslungenes Exempel war der Chatbot Tay, der im Frühjahr 2016 lernen sollte, wie sich Menschen im Internet unterhalten. Keine 24 Stunden später musste Microsoft ihn wieder vom Netz nehmen, weil er zum Rassisten geworden war, der den Holocaust leugnete und Hitler lobte. Das Programm hatte offensichtlich von den falschen Leuten gelernt und agierte ohne Hemmungen, weil keine Regeln eingebaut waren.

Doch wenn Maschinen immer eigenständiger handeln, dann müssen sie zwischen Falsch und Richtig, Gut und Böse unterscheiden können. Soll ein Roboter Familienangehörige verständigen, wenn ein Senior seine Medikamente nicht nimmt, oder ist dessen Selbstbestimmung höher zu bewerten? Soll ein selbstfahrendes Auto vor einem Igel bremsen, wenn ein nachfolgendes Fahrzeug einen Auffahrunfall verursachen könnte? Oder schlimmer: Wie entscheidet so ein Auto, wenn es nur die Wahl hat, einen unvorsichtigen Passanten zu überfahren oder mitsamt seinem Fahrgast mit einem Baum zu kollidieren?

Einem Menschen würde man eine falsche Entscheidung, die er in Sekundenbruchteilen trifft, wohl verzeihen – aber einer Künstlichen Intelligenz? Muss sie nicht alle Eventualitäten im Voraus berücksichtigen? In der Praxis ist das allerdings oft gar nicht möglich. Abgesehen davon, dass auch eine KI nicht alles berechnen kann – wie etwa die exakten Reibungs-

kräfte auf einer glatten Straße –, so ist es in Deutschland unzulässig, potenzielle Opfer gegeneinander aufzurechnen. Lässt sich daher ein Unfall mit Personenschaden nicht vermeiden, würde die Maschine wohl scharf bremsen, ein Schleudern in Kauf nehmen und der Physik ihren Lauf lassen müssen.

Vielfach wird auch diskutiert, ob man nicht die »Drei Gesetze der Robotik«, die der Science-Fiction-Autor Isaac Asimov bereits 1942 formulierte, als fundamentale Regeln in die Maschinen integrieren sollte. So besagt das erste Gesetz: Ein Roboter darf keinen Menschen verletzen oder durch Untätigkeit zu Schaden kommen lassen. Das zweite: Ein Roboter muss den Befehlen eines Menschen gehorchen, es sei denn, sie stehen im Widerspruch zum Ersten Gesetz. Und das dritte: Ein Roboter muss seine eigene Existenz schützen, solange dies nicht mit dem Ersten oder Zweiten Gesetz kollidiert.

Doch dabei ergeben sich etliche Probleme, etwa die Frage, was ein »Schaden« ist. Fällt darunter zum Beispiel auch ein finanzieller Verlust oder Rufschädigung? Je nach kultureller Umgebung werden die Prioritäten vielleicht unterschiedlich gesetzt. Außerdem hat Asimov selbst Fälle konstruiert, in denen seine Regeln sogar verletzt werden müssen, weil sich logische Konflikte ergeben oder weil ein Roboter eine Situation anders einschätzt als ein Mensch.

Letztlich wollen Maschinenethiker KI-Systeme konstruieren, die so verlässlich agieren wie ein Geschirrspüler oder ein Putz-Roboter – diesen Geräten vertrauen wir so sehr, dass wir sogar die Wohnung verlassen, während sie arbeiten. Bei einfachen Systemen könnte man etwa den Käufer über die Moral seiner Maschine entscheiden lassen. So haben Forscher in der Schweiz einen Putzroboter entwickelt, der sich so einstellen lässt, dass er zwar Spinnen einsaugt, aber Marienkäfer nicht – oder beide nicht, je nachdem, was der Kunde möchte. Genauso könnte man bei Drohnen festlegen, dass sie zwar

Tiere und Pflanzen fotografieren dürfen, aber keine Menschen. Und bei Autos, dass sie bei Schnecken und Insekten nicht bremsen, aber bei Igel und Kröten, wenn es die Verkehrssituation erlaubt, und auf jeden Fall bei größeren Tieren wie Rehen und Wildschweinen.

Vielfach wird man auch strikte Regeln einpflanzen müssen, etwa »Menschenleben sind wichtiger als Sachschäden«, oder bei einem Chatbot, dass lobende Aussagen über Hitler tabu sind, dass er nicht absichtlich die Unwahrheit sagt, dass er nicht so tut, als ob er ein Mensch sei, und dass er die Probleme des Benutzers ernst nimmt. So soll er angemessen reagieren, wenn sein Gegenüber depressiv ist – und betonen, dass er nur eine Maschine sei und ein psychologisches Gespräch mit einem Menschen empfehle.

Zugleich müssen die Maschinen lernen, was ein sozial akzeptiertes Verhalten ist. So werden Fahrzeuge, die auf der Autobahn hinter anderen herschleichen, von den übrigen Verkehrsteilnehmern ebenso negativ bewertet wie solche, die ständig die Spuren wechseln. Selbsttätig fahrende Autos sollten also ab und zu andere überholen, aber nicht übertrieben rücksichtslos fahren – ganz wie vernünftige menschliche Autofahrer.

Falls eine KI vor unbekannten Problemen steht, kann man ihr auch Datenbanken zur Verfügung stellen, in denen sie Lösungen für vergleichbare Situationen finden kann. Außerdem versuchen Forscher, KI-Systeme eigene Schlüsse aus Beobachtungen ziehen zu lassen – etwa indem sie aus geeigneten Filmen oder Büchern lernen, wie sich Menschen üblicherweise verhalten. In virtuellen Welten werden dann schwierige Situationen und Dilemmata simuliert, und die KI wird über ein Punktesystem belohnt, wenn sie moralisch richtig entschieden hat. Solche Forschungen stehen allerdings noch ganz am Anfang.



Gibt es etwas, das Computer, Roboter und andere Systeme mit Künstlicher Intelligenz niemals erreichen werden? Man soll zwar nie »nie« sagen, aber zwei Hürden scheinen in der Tat praktisch unüberwindlich zu sein – zumindest für Maschinen, wie sie heute erdacht werden. Das eine ist der Unterschied zwischen der Biologie und elektronischen Schaltungen, das andere sind die Erfahrungen eines menschlichen Lebens.

Viele Wissenschaftler bezeichnen unser Gehirn mit seinen 86 Milliarden Nervenzellen, den Neuronen, und mehreren Hundert Billionen Verknüpfungen, den Synapsen, als »das komplexeste Objekt des bekannten Universums«. Das hat seinen Grund: Denn anders als ein klassischer Computer speichert es nicht digitale Werte, also Nullen und Einsen, sondern Hunderte von Millionen oder Milliarden von analogen Mustern: Bilder und Abläufe, Begriffe und Bedeutungen, Töne und Klänge, Gerüche, Gefühle, Bewertungen. Vergleicht man die Anzahl gespeicherter Muster und ihre gleichzeitige Verarbeitung, dann liegt die Leistungsfähigkeit eines Gehirns etwa in der Größenordnung heutiger Superrechner. Nur dass ein einziger dieser Superrechner so viel Energie verbraucht wie eine mittelgroße Stadt. Das menschliche Gehirn hingegen begnügt sich mit einem Millionstel davon: Gerade einmal 20 Watt – so viel wie eine kleine Lampe.

Außerdem funktioniert unser lernfähiges Denkorgan ohne Software, ohne zentrale Steuerung und ohne Betriebssystem, und es ist extrem fehlertolerant: Obwohl jeden Tag etwa

100 000 Neuronen verloren gehen, lassen seine kognitiven Fähigkeiten über Jahrzehnte hinweg kaum nach. Das Gehirn kann mit verlorenen Ressourcen ebenso gut umgehen wie mit unpräzisen Informationen. Und noch etwas unterscheidet biologische Systeme von elektronischen: Sie sind in der Lage, sich bis zu einem gewissen Grad selbst zu reparieren. Zellen wachsen neu, das Immunsystem bekämpft Eindringlinge, Abfall wird zerlegt und abtransportiert. Biologische Systeme wachsen, sie passen sich flexibel an neue Herausforderungen an – und sie vermehren sich, bekommen Kinder.

KI-Systeme können zwar besser als unser Gehirn enorme Datenmengen durchforsten, sie können schneller und präziser Muster in Bildern, Sprache und Texten lernen und erkennen. Sie können uns auf vielen Feldern helfen, können Wissenszusammenhänge herstellen und Informationsbausteine kombinieren, sind in gewissem Maße sogar kreativ, aber in den Punkten Energieverbrauch, Flexibilität, Fehlertoleranz und Reparaturfähigkeit sind sie von biologischen Systemen sehr weit entfernt.

Auch das fehlende Verständnis unserer Welt wird sich nicht einfach beheben lassen. Ein gutes Beispiel sind sogenannte Winograd-Sätze. Der Informatiker Terry Winograd hatte in den 1970er-Jahren viele Sätze entwickelt, die sich nur in einem Begriff unterscheiden und an denen sich zeigt, wie viel Hintergrundwissen nötig ist, um sie richtig zu interpretieren. »Die Behördenvertreter verboten den Demonstranten, sich zu versammeln, weil sie Gewalt befürchteten« beziehungsweise »Die Behördenvertreter verboten den Demonstranten, sich zu versammeln, weil sie Gewalt befürworteten« ist so ein Winograd-Satzpaar. Für einen Computer oder ein anderes der heutigen KI-Systeme ist es extrem kompliziert, herauszufinden, auf wen sich das »befürworten« beziehungsweise »befürchten« bezieht – ein Mensch erkennt das sofort.

Oder andere Beispiele: »Der Ball durchbrach den Tisch – er war aus Styropor« und »Der Ball durchbrach den Tisch – er war aus Stahl« sowie »Der Pokal passte nicht in den Rucksack, er war zu groß«. Worauf bezieht sich jeweils das »er«? Auch dafür braucht man einiges an Alltagserfahrung. Natürlich lässt sich manches in Datenbanken hinterlegen, aber oft wissen wir gar nicht, was wir alles wissen. Nehmen wir nur ein Foto von Leuten auf einer Party: Wir Menschen sind sehr gut darin, die Szene zu interpretieren und vielleicht sogar zu erkennen, wer wen verliebt anschaut – Computer können das nicht.

Was würde wohl ein KI-System auf diese Frage antworten: »Es sitzen 15 Vögel auf einer Stromleitung. Der Jäger erschießt einen. Wie viele bleiben sitzen?« Vermutlich »14«, denn die richtige Antwort erfordert einiges an Zusatzwissen: dass ein Schuss knallt und dass Vögel bei einem lauten Knall erschreckt auffliegen und sich anderswo einen ruhigeren Ort suchen. Das Interessante an Menschen ist, dass in diesem Fall auch jemand, der noch nie gesehen hat, wie ein Jäger auf Vögel schießt, durch Nachdenken und Intuition auf die richtige Lösung kommen kann. Doch woher soll eine Maschine diesen »gesunden Menschenverstand« haben?

Vieles von dem, was wir im Alltag nutzen, findet man nicht im Internet. Ebenso ergeht es Robotern, wenn sie Dinge suchen, die sie für ihre Aktionen brauchen – man denke nur an die Reibungskräfte beim Anheben von Teetassen. Die stehen auf keiner Webseite. Wenn Maschinen so etwas lernen wollen, müssen sie uns im Alltag begleiten, sozusagen mit uns aufwachsen, wie dies kleine Kinder tun. Und selbst dann wird ihnen vieles verborgen bleiben: Sie wissen nicht, wie es ist, zu essen und zu trinken, zu schlafen und zu träumen, zu wachsen und zu altern, sich zu verlieben und Kinder zu bekommen – ihre Maschinen-Gefühle werden immer andere sein als die durch Hormone verursachten Menschen-Gefühle.

Das ist auch der Grund, warum es für Maschinen so schwierig ist, den Turing-Test zu bestehen. Schon 1950 hatte sich der Mathematiker Alan Turing die Frage gestellt: »Können Maschinen denken?« Um sie zu beantworten, schlug er einen Test vor: Führt ein Mensch eine Unterhaltung mit einer Maschine, die er nicht sehen kann, und er ist sich danach unsicher, ob sein Gegenüber nicht doch ein Mensch ist, so hat die Maschine den Test bestanden. 40 Jahre später lobte der Soziologe Hugh Gene Loebner einen Preis für diejenige Software aus, die in einem 25 Minuten dauernden Wortwechsel den Turing-Test erfolgreich absolviert.

Doch bis heute musste das Preisgeld von 100 000 Dollar noch nie ausbezahlt werden. Zwar werden die Dialog-Programme immer besser, aber manchmal weichen sie ungeschickt aus, verwenden Textbausteine mehrfach oder wissen auf einfache Alltagsfragen keine sinnvolle Antwort. Die Liste der digitalen Bewerber, die immerhin eine »Bronzemedaille« erhielten, führt derzeit die Chatbot-Lady Mitsuku an. Inzwischen ist auch klar, worin das Problem besteht: Der Turing-Test findet nicht so sehr heraus, ob Maschinen denken, sondern er misst ihre Menschenähnlichkeit – und da werden sie noch lange nicht überzeugen können.

Letzten Endes muss sich jedes Wesen, das mit der Welt umzugehen hat, in ihr selbst entwickeln, mit einem Körper, Sinnesorganen und praktischer Intelligenz. Darin aber sind wir Menschen kaum zu schlagen, denn im Lauf von Millionen Jahren der Evolution hat sich unser Gehirn als so anpassungs- und lernfähig erwiesen, dass es auch mit ganz neuen Herausforderungen umgehen kann – sogar wenn es sie selbst entwickelt hat, wie die smarten Maschinen.