

Zeitschrift: Schweizer Monat : die Autorenzeitschrift für Politik, Wirtschaft und Kultur
Band: 96 (2016)
Heft: 1036

Artikel: Schau, wie die Maschine denkt
Autor: Pines, Sarah
DOI: <https://doi.org/10.5169/seals-736308>

Nutzungsbedingungen

Die ETH-Bibliothek ist die Anbieterin der digitalisierten Zeitschriften. Sie besitzt keine Urheberrechte an den Zeitschriften und ist nicht verantwortlich für deren Inhalte. Die Rechte liegen in der Regel bei den Herausgebern beziehungsweise den externen Rechteinhabern. [Siehe Rechtliche Hinweise.](#)

Conditions d'utilisation

L'ETH Library est le fournisseur des revues numérisées. Elle ne détient aucun droit d'auteur sur les revues et n'est pas responsable de leur contenu. En règle générale, les droits sont détenus par les éditeurs ou les détenteurs de droits externes. [Voir Informations légales.](#)

Terms of use

The ETH Library is the provider of the digitised journals. It does not own any copyrights to the journals and is not responsible for their content. The rights usually lie with the publishers or the external rights holders. [See Legal notice.](#)

Download PDF: 29.04.2025

ETH-Bibliothek Zürich, E-Periodica, <https://www.e-periodica.ch>

Schau, wie die Maschine denkt

Was ist künstliche Intelligenz? Und warum kann sie einem durchaus Angst einjagen? Ein Laborbesuch im Silicon Valley, wo neue Maschinenbabys sorgfältig gezüchtet werden.

von Sarah Pines

Der ruhig leuchtende Touchscreen des Tesla zeigt hinter verdunkelten Scheiben den Selbstfahrmodus an. Auf der 90 Stundenkilometer schnellen Fahrt über den Highway liegen meine Hände im Schoss; mit zaghaft kleinem Ruckeln reguliert der Wagen Abstände zu anderen Autos und Leitplanken. Die Route: Arastradero-Naturpark, Eukalyptusbaumgeruch beim kurzen Herunterfahren der Fenster, dann Junipero Serra Freeway, Sand Hill Road, irgendwann die Kurve zum Palm Drive, der mit Palmen flankierten Auffahrt der Stanford University. Das Lenkrad schwenkt herum, der Tesla parkt sich am «Oval», der schlaufenförmigen Parkstruktur vor dem Hauptgebäude; ich steige aus. An einer Ecke mit Oleanerbaum winkt ein junger Mann in Shorts, Käppi und mit Ray Ban eine Gruppe Studenten über die Strasse. Direkt vor ihnen kommt ein weisses Auto in Marshmallowform mit winzigem Surren zum Stehen. Das Radarsystem auf dem gewölbten Dach ähnelt einem in eine Zahnarztzange geklemmten Wasserkocher, darin erfasst ein Laser die Umgebung – Strasse, Verkehrsschilder, Fussgänger – und sorgt dafür, dass der Marshmallow Hindernissen ausweicht und rechtzeitig bremst. Der junge Mann mit Käppi markiert etwas auf einem iPad-Touchscreen. Ein anderer Mann in Anzughose und Shirt mit Uni-Logo steigt aus dem Wagen, High-Five! Der selbstfahrende Google Car hat eine weitere Testfahrt bestanden.

Sarah Pines

ist Kulturjournalistin und lebt in Palo Alto, Kalifornien.

Zweck einer AI (engl. Artificial Intelligence, künstliche Intelligenz) wie des Google Car? «Sichere Strassen», antwortet der Mann mit Ray Ban, wischt kurz übers iPad, Headsetblinken, eingehender Anruf. «Warte eine Sekunde, ja?» Dann weiter: den Verkehr entstopfen, zeternde Autofahrer, Parkstress – all das in Zukunft abschaffen. Sauber fliessender, fahrerloser Verkehr. Das ist im Silicon Valley Untertraum einer Grossvision – einer Welt voller künstlicher Intelligenz.

Das Träumen begann, als 1963 der Erfinder des AI-Begriffs John McCarthy in Stanford das erste Artificial Intelligence Laboratory der Vereinigten Staaten gründete. Das Labor befindet sich



heute in der linken Flanke des ockerfarbenen Uni-Hauptgebäudes und brachte insgesamt 17 Turing-Medaillen-Gewinner hervor, dieses Nobelpreises der Computerwissenschaften. Absolventen, Studierende und Lehrende arbeiten eng mit den Firmen des Valley zusammen – Google, Facebook, Microsoft –, die die Entwicklung künstlicher Intelligenz fördern, weltweit anführen und für ihre Forschung wiederum die besten Neuronenforscher, Programmierer, Computertechnologen aus Stanford anheuern. Direktorin des Artificial Intelligence Laboratory ist heute Professorin Fei-Fei Li, eine zierliche und jung wirkende Physikerin und Computerwissenschaftlerin, die schon als Studentin internationale Auszeichnungen gewann. In der derzeitigen Forschung gehe es grundsätzlich darum, sagt Li, dass Computer Informationen analysierten und interpretieren lernten wie Menschen. Die Grundeinheit dabei ist der Algorithmus, eine in Zahlen kodierte Anweisung. Forschungsziel ist es, Algorithmen zu schaffen, die die menschliche Neuronenstruktur perfekt simulieren, also «denken» können. Das gelingt heute bereits in Ansätzen.

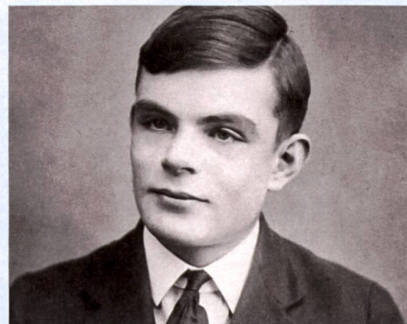
Computerprogramme können einfache Umgebungen und Sachverhalte erfassen, wie: «Das auf dem Tisch ist eine Tasse, darin ist eingetrockneter Kaffeesatz; aus dieser Tasse wurde also getrunken.» Sie können schlichte Befehle verstehen und ausführen, aber auch in Sekundenschnelle Abermillionen Daten erfassen und hochkomplexe Aufgaben lösen. Ein paar Forschungshöhepunkte: 1997 besiegte der IBM-Computer Deep Blue den Schachweltmeister Gary Kasparow, 2011 gewann Watson, ebenfalls von IBM, das TV-Quiz «Jeopardy», diesen März schlug Google DeepMinds AlphaGo den Profi Lee Sedol im Brettspiel Go. Apple entwarf Siri, Google entwickelte 2012 eine Gesichtserkennungssoftware (für Katzen), dann Spot, den trappenden, schnauzenlosen Roboterhund, Spionagetier für Kriegsgebiete, schliesslich Atlas, den bisher fortschrittlichsten Roboter. Kurz vor seinem Tod in Stanford 2011 bezeichnete McCarthy solch erste robotische Aufgabenbewältigungen als «Drosophila der künstlichen Intelligenz» – primitiv, aber wie die Fruchtfliege als Grundlage der Genetikforschung am Menschen unerlässlich auf dem Weg zum eigentlichen Forschungsziel: der intellektuellen Selbständigkeit von Maschinen.

Denken als Abfolge von Algorithmen

Die Idee von der intellektuellen Gleichwertigkeit von Mensch und Maschine, der sogenannten «Singularität», geht zurück auf Alan Turing. Während des Zweiten Weltkriegs baute der britische Mathematiker im Spionagezentrum Bletchley Park im Auftrag von Winston Churchill eine Maschine, die den Funkcode der Nationalsozialisten knackte und den Sieg der Alliierten beschleunigte. Nach dem Krieg lehrte Turing Mathematik an der Manchester University und begann, sich für die Frage nach der Intelligenz von Maschinen zu interessieren: Kann eine Maschine Denken so weit erlernen, dass es von menschlichem Denken ununterscheidbar wird? Turings Antwort: ja, wenn es gelingt, das menschliche Denken in Serien von Algorithmen zu übersetzen und Maschinen mit

Alan Turing

1912–1954



Britischer Logiker, Mathematiker, Kryptoanalytiker und bis heute einer der bedeutendsten Denker der theoretischen Informatik. Bereits als Jugendlicher begann Turing sich mit der Frage auseinanderzusetzen, was Denken sei und wie es funktioniert.

1937 veröffentlichte er seinen Aufsatz «On Computable Numbers, With an Application to the Entscheidungsproblem». Darin beschrieb er eine Maschine, die als «Turingmaschine» in die Geschichte eingehen sollte und die, theoretisch, fähig war, jede berechenbare Zahlenfolge zu berechnen («to compute»). Dafür galt – nach Vorbild des menschlichen Denkens, dessen Rechenleistung begrenzt ist – für den Ausdruck «berechenbar», dass in absehbarer Zeit ein Resultat erfolgt. Eine Maschine, die ewig rechnet, ist unbrauchbar. In einer begrenzten Anzahl von Berechnungen kommt eine Maschine nur zu einem Resultat, wenn sie durch einen Algorithmus programmiert ist.

Während des Zweiten Weltkriegs arbeitete er als Kryptoanalytiker für den britischen Geheimdienst und war beteiligt an der Entschlüsselung des Codes der deutschen Enigmamaschine, die die Kommunikation der Nationalsozialisten verschlüsselte. 1950 erscheint sein Artikel «Computing Machinery and Intelligence», in dem er zur Debatte stellt, ob Maschinen denken können. Er sieht Fortschritte in der Programmierung und Speicherkapazität von digitalen Computern voraus, so dass 50 Jahre später, also im Jahre 2000, nicht mehr erkennbar sei, ob man sich mit einem Menschen oder einer Maschine unterhalte. Diese als «Turing-Test» bekannte Unentscheidbarkeit gilt noch immer als ein Kriterium für «künstliche Intelligenz». Die Programmierung einer dazu fähigen Maschine dahin sieht er in zwei Teilen: einen rudimentären Verstand entwickeln – dem eines Kindes ähnlich – und diesen dann aus- bzw. weiterbilden. Diese Arbeitsweise hat sich bis heute in der Informatik etabliert. (SJ)

diesen Algorithmen zu programmieren. Wenn Mensch und Maschine dann in ihrer Intelligenz ununterscheidbar werden, ist der berühmte «Turing-Test» bestanden, Singularität erreicht.

Ava aus «Ex Machina», «Darth Vader», «Robocop», «Terminator» – die Klischeeform künstlicher, singulärer Intelligenz ist der «Cyborg», eine Vorstellung, die in den 1960er Jahren im Sonoma County nördlich des Silicon Valley entstand. In einem wissenschaftlichen Artikel der Zeitschrift «Astronautics» von 1960 verwendete der Forscher Manfred Clynes, der, inzwischen 90jährig, immer noch in Sonoma lebt, erstmalig den Begriff Cyborg, um einen Maschinen-Mensch-Hybriden zu beschreiben. Die Menschheit werde, so seine Idee, das All erforschen, eine völlig andere Umgebung als die Erde, mit neuen Lichtverhältnissen, Luftströmen und Geschwindigkeiten. Wenn Mensch und Maschine aber verschmelzen, so Clynes weiter, liessen sich widerstandsfähigere Körper schaffen, und der Geist könnte sich ganz der Erfahrung der neuen Welt hingeben. Ähnlich sieht das Ray Kurzweil, Googles Forschungsdirektor. Er prophezeit Singularität für das Jahr 2029. Kurzweil aber geht einen Schritt weiter. Er glaubt an eine Zukunft, in der Maschinen den Menschen an Intelligenz übertreffen, die sogenannte «Intelligenzexplosion». Der Begriff stammt vom Kryptologen I. J. Good, einst Freund und Kollege von Alan Turing in Bletchley Park. 1960 hatte der Psychologe Franz Rosenblatt einen raumgrossen Computer namens Perceptron geschaffen, der einfache visuelle Strukturen erfasste und verstand. Good, der damals am Massachusetts Institute of Technology lehrte, inspirierte das zum Aufsatz «Speculations Concerning the First Ultraintelligent Machine». Darin benannte Good erstmals die Idee künftiger superintelligenter Maschinen. «Das Überleben des Menschen hängt von der baldigen Konstruktion einer ultraintelligenten Maschine ab», schreibt Good zu Beginn des Textes, dann weiter:

«Eine ultraintelligente Maschine übertrifft die intellektuellen Aktivitäten auch des cleversten Menschen bei weitem. Da die Herstellung einer solchen Maschine selbst eine intellektuelle Aktivität darstellt, könnte eine ultraintelligente Maschine noch bessere Maschinen herstellen, es käme zu einer Intelligenzexplosion, die Intelligenz des Menschen würde weit übertreffen. Also ist die erste ultraintelligente Maschine die letzte Erfindung, die der Mensch je machen muss, vorausgesetzt, dass die Maschine gefügig genug ist und uns sagt, wie man sie unter Kontrolle halten kann.»

Good war überzeugt, dass die politischen Geschehnisse seiner Zeit – der Kalte Krieg, das nukleare Wettrüsten zwischen den USA und Russland, die Kubakrise – den Menschen an den Rand der Auslöschung bringen würden. Nachdem, so der Gedanke, im Zweiten Weltkrieg eine intelligente Maschine schon einmal den Untergang der Welt verhindert hatte, brauchte die Welt nun eine neue, bessere künstliche Intelligenz. Bereits Good baute in seinen Text Bemerkungen der Vorsicht ein. In Science Fiction, so schrieb er unter anderem, werde maschineller Ungehorsam oft themati-

siert – er verstehe nicht, warum dies sonst so selten getan werde. «Vorausgesetzt, dass die Maschine gefügig ist», schrieb Good bezeichnenderweise. Falls.

Das Ende der Menschheit

Die Angst vor tödlichen Maschinen hat die Menschen nie verlassen. Selbst im Silicon Valley nicht. Heute warnen dort Stephen Hawking und Steve Wozniak, Mitgründer von Apple, vor dem Ende der Menschheit, dem Dräuen der Maschine. Etwas pragmatischer sieht es Elon Musk, derzeitiger Visionär des Silicon Valley, Erfinder des Tesla und Gründer von SpaceX, der ersten Firma, die Raumschiffe billiger, schneller und besser entwickeln kann als die Nasa. In 500 Millionen Jahren wird die Erde laut Musk überhitzt, aufgeschwollen, pflanzen- und wasserlos der Zerstörung durch die Sonne harren. Schon lange vorher werden der Erde die Ressourcen ausgehen, die Menschheit wird entweder aussterben oder sich im Weltraum andere Planeten erschliessen. Da wir, so Musk weiter, noch Zeit brauchen, um das Universum ausreichend zu verstehen, kaufen Elektroautos, Solartechnik, aber auch künstliche Intelligenz uns die nötige Zeit, die globale Erwärmung zu verhindern, die Raumfahrt zu vervollkommen und den Mars zu bevölkern. Allerdings sei bei der Vervollkommenung künstlicher Intelligenz Vorsicht geboten, warnt Musk in der Tradition I. J. Goods. Es brauche Kontrollinstanzen; Forschungsgelder und Spenden sollten auch der Förderung der Sicherheit von Computertechnologien dienen.

Noch tapsen die denkenden Maschinen eher tollpatschig als furchteinflössend durch die Welt. Videos der amerikanischen DARPA Robotics Challenge von 2015 zeigen ein Konglomerat von Robotern, die mit der Geschwindigkeit rückenkranker Pferde oder schlurfender Greise Geräte fahren, über unwegsames Gelände laufen, Türen öffnen, Sachen aufheben und so fort. Der virtuelle weibliche Twitter-Teenager Tay (@tayandyou) von Microsoft verfiel diesen Frühling, provoziert und geschult von Tweets der Nutzergemeinde, in Hitler-verherrlichende, sexistische Hasstiraden – und darf jetzt nur noch mit ausgewählten Menschen sprechen. Letzten Sommer verpasste ein Verkehrspolizist einem Google Car einen Strafzettel für zu langsames Fahren, ein anderer schrammte einen Bus, als er einem Sandsack auf einer Baustelle auswich. (Frage an die Juristen unter unseren Lesern: Wer war schuld, der auf der Rückbank dösende Mitfahrer oder das Auto? Kann man eine Software verklagen?) Woanders gehen Kunden virtuelle Callcentermitarbeiter auf die Nerven. Kurzum: gegenwärtigen «komplexen» Robotersystemen gelingen höchstens ein paar trappelnde Schritte, ein tapsiges Türöffnen, eine Schachrunde. Warum also diese Angst vor Sternenkriege auslösenden, technophilen Utopien, vor kosmonautischem Grössenwahn?

Kein Sinn für Ironie und Metaphern

Die geläufige Antwort ist eine ganz grundsätzliche: es sei schwer vorstellbar, dass künstliche Intelligenz einmal über dieselbe emotionale Komplexität verfügen, zwischen Gut und Böse

unterscheiden könne wie ein Mensch. Könne eine Maschine je Ironie, Sarkasmus, Zynismus verstehen, würde sie Befehle nicht stets wörtlich nehmen und stur ausführen? Ein Beispiel: ich schmeisse mich eilig auf den Rücksitz meines selbststeuernden Autos, rufe: «So schnell wie möglich zum Flughafen!» Die Befürchtung: das Auto wird den Befehl ausführen, vielleicht auf der rasend schnellen Fahrt rote Ampeln, zur Seite springende Fussgänger und den schwindlig erbrechenden, Haltebefehle rufenden Mitfahrer ignorieren. Was ist mit selbstschiessenden Nuklearwaffensystemen? Wer stoppt eine Maschine, die, intelligenter als der Mensch, eine noch viel intelligentere Maschine baut, diese so optimiert, dass sie sich dem Abstellen widersetzt? Also unbeherrschbar wird? Wird die Maschine den Menschen ersetzen wie einst der Mensch den Neandertaler?

Im Februar präsentierte Googles Tochterfirma Boston Dynamics in einem YouTube-Video eine optimierte Version des Roboters Atlas. Mit leisem Surren stapft Atlas durch einen verschneiten Wald, knickt an unebenen Stellen mit dem Knie leicht ein, fängt sich, stapft weiter. Dann stösst ein Techniker Atlas teilweise mit einem Hockeyschläger um. Er schlägt vornüber, die Beine vor dem Bauch gekreuzt wie ein erschreckt eingefalteter Marienkäfer; auch beim DARPA-Wettlauf erinnern die Roboter antreibenden Menschen an Käfer pieksende Tierquäler. Beim Beobachten ein nervöses Gedankenflattern: wird Atlas, werden die Robotergerisse der DARPA sich an den Angreifern rächen? Wenn nicht jetzt gleich, dann eines Tages?

Warnungen tut die Forschungswelt des Silicon Valley meist als anthropozentrische Existenzangst ab. Sei es nicht herrlich, wenn kluge Maschinen dem Menschen beistünden, intelligenter

und fitter zu werden, vielleicht gar ewig zu leben? So sieht das Google-Forscher Kurzweil. Das Ziel künstlicher Intelligenz, so der Tenor im Silicon Valley, sei nicht unbedingt das Bestehen von Turing-Tests oder verquere Antworten auf das Descartes'sche Leib-Seele-Problem. Man denkt viel pragmatischer. Bisher, so das Gerücht, ist das US-Militär Hauptabnehmer der von Google fabrizierten Roboter: Gladiatorenhaft über Geröllfelder trabende Androide, pillenförmige, gelbe Roboterfische, die ab 2017 die abgerichteten Seehunde und Delphine der US-Navy ersetzen. Netzwerke vogelschwarmartig operierender Marschflugkörper, Überwachungsroboter in Taubengestalt oder winziger Maikäferform, reglos auf Antennenmasten, Dächern, Mauern sitzend. Auch andere Kunden aber sollen profitieren. Künstliche Intelligenz sei im Mainstream angekommen, verkündete am 28. März ein Aufsatz der «MIT Technology Review», sei für Versicherungs- und Finanzfragen, Autoherstellung und Ölgewinnung einsetzbar. Anfang Januar gab Facebook-Chef Mark Zuckerberg seinen Neujahrsvorsatz bekannt: Er wolle einen Jarvis bauen, einen omnipräsenten Softwaremembran-Butler wie der des Superhelden Iron Man, der die Belange seines Alltags regeln solle.

Man habe nichts im Sinn, so sagte es bereits der Mann mit Käppi vor der Stanford University, als das Leben einfacher und effektiver zu gestalten: Kriege mit maschinellen und nicht menschlichen Todesopfern zu führen, Roboter zu erfinden, die dem Menschen assistieren. In Fragen der Medizin, der Rente, des Transportwesens oder sozialer Ungleichheit mit besseren Lösungen aufzuwarten, als der Mensch sie sich auszudenken vermöchte.

«Move fast», so lautet Zuckerbergs berühmter Satz, «and break things.» ◀

Anzeige

Familie Zahner | 8467 Truttikon

052 317 19 49 | www.zahner.biz | zahner@swissworld.com

Klassische Flaschengärung in unserem Keller.

Fr. 20.—

Truttiker
Schaumwein

Blanc de Pinot Blanc brut