

»Wir sind viel berechenbarer, als wir denken«

Katharina Zweig ist Forscherin und Informatik-Professorin und beschäftigt sich unter anderem mit der Analyse komplexer Netzwerke. Wir sprachen mit ihr darüber, was Algorithmen können – und was nicht. Und was das für unsere Möglichkeiten und menschlichen Pflichten in der Zukunft bedeutet.

Interview: Maja Beckers, Greta Lührs

IM HAMBURGER STADTTEIL HOHELUFTE kennt Katharina Zweig sich aus, denn sie ist ganz in der Nähe aufgewachsen. Zu unserem Treffen in einem Café kommt sie mit einem Buch über Verschwörungstheorien in der Hand – ein Thema, auf das wir im Laufe des Gesprächs noch zurückkommen werden. Die in Kaiserslautern lehrende Professorin ist in ihrer Heimatstadt, um einen Vortrag zu halten, und hat sich die Zeit für ein Gespräch mit uns genommen, weil ihr die Verbindung von Informatik und Philosophie am Herzen liegt. Ihrer Meinung nach begegnen sich beide Disziplinen mit zu viel Argwohn. »Wir brauchen aber Leute, die beides zusammen denken können«, sagt sie. Zwar müsse nicht jeder Informatiker Philosoph werden und andersherum, doch mehr Annäherung hält sie für wichtig. Zweig vereint in ihrem Lebenslauf selbst verschiedene Interessen: Sie studierte zunächst Biochemie, dann kam Informatik hinzu, und schließlich »verliebte« sie sich in die theoretische Informatik. Die führte sie wiederum zu Grundsatzfragen – wie: Was lässt sich überhaupt berechnen? Nach dem Gespräch sagt Katharina Zweig, sie hätte gern noch mehr philosophiert. Ob wir das beim nächsten Mal nachholen könnten?

HOHE LUFT: Professorin Zweig, Algorithmen sind für viele ein Schreckgespenst, sie haben etwas Mystisches, das auch Angst machen kann. Können Sie das nachvollziehen?

KATHARINA ZWIG: Das liegt daran, dass es eine große Verwirrung gibt. Es gibt die Algorithmen, vor denen keiner Angst hat. Das Navigationsgerät berechnet einfach verlässlich den kürzesten Weg von A nach B, und davor fürchtet sich niemand. Eigentlich ist es nur ein ganz kleiner Teil, vor dem wir Respekt haben sollten. Das sind algorithmische Entscheidungssysteme, die aus Daten aus der Vergangenheit etwas darüber lernen, wie Menschen sich verhalten haben, und daraufhin Entscheidungen über sie treffen. Wir werden in vielen Bereichen unseres Lebens schon durch Systeme klassifiziert, und in den meisten Fällen ist das auch in Ordnung, wenn das Regelwerk einsichtig ist. Jeder, der Auto fährt, ist zum Beispiel in einer Schadensfreiheitsklasse. Als 18-Jährige ärgert man sich, dass man so viel bezahlen muss, aber es ist nachvollziehbar. Wir wissen, dass 18-Jährige mit ihrer mangelnden Erfahrung eben nicht besonders gut fahren. Das ist ein Klassifizierungssystem, gegen das ich noch nie Protest gehört habe. Aber wenn man nun

diese Methoden automatisiert anwendet, um herauszufinden, ob ein Krimineller vermutlich eher rückfällig wird oder nicht – spätestens da fängt es an, bedenklich zu werden. Denn erstens steht viel auf dem Spiel, vielleicht ein oder zwei Jahre länger im Gefängnis, und zweitens bin ich diesem Algorithmus ausgeliefert. Ich kann nicht, wie bei der Autoversicherung, einfach zu einem anderen Anbieter gehen. Das heißt, in dem Moment, wo man mithilfe von Algorithmen aus Daten lernt, wie sich Menschen verhalten haben, und daraus extrapoliert, wie es in der Zukunft sein könnte, wird es unangenehm, und da müssen wir hingucken.

Wo werden solche algorithmischen Entscheidungssysteme angewandt?

Das ist eines der Probleme, wir wissen nicht genau, wo sie wie verwendet werden. Wir wissen nur von wenigen Fällen, in Polen wird es etwa verwendet, um Arbeitslose zu klassifizieren, in den USA, um genau solche Rückfallquoten zu berechnen. Der China-Citizen-Score wird wahrscheinlich auch eine lernende Komponente mit beinhalten, und in den Snowden-Dokumenten wird ein System namens Skynet erklärt, das von verurteilten und verdächtigten Terroristen lernen soll, welche Eigenschaften terroristische Kuriere haben. Es wird jetzt diskutiert, ob man Systeme, die Terroristen identifizieren und dann auch gleich schießen dürften, verbannen kann – aber nicht alle Staaten stimmen dem zu. Neben diesen realen Beispielen gibt es viele Studien, die alle hinreichend gruselig sind.

Inwiefern?

Sie behaupten, Algorithmen könnten Depressive auf Instagram finden oder Homosexuelle oder Kriminelle per Foto erkennen – das sind Geschichten, die man besser nicht so einfach glauben sollte. Bei einer dieser Studien gab es zum Beispiel eine Reihe von Verbrechern, deren Fotos man dem System gegeben hat, und eine Reihe von Nicht-Verbrechern. Das System sollte sie auseinanderhalten. Leider waren die Kriminellen-Bilder alle Mugshots, das heißt im Gefängnis

aufgenommen im T-Shirt, und die Bilder der Nicht-Kriminellen waren Bewerbungsfotos in Anzug und Krawatte. Also ist es sehr wahrscheinlich, dass das System einfach »verstanden« hat: Du trägst eine Krawatte, du bist ein guter Mensch – und natürlich nicht wirklich Kriminelle von Nicht-Kriminellen unterscheiden kann.

Wie kann das passieren?

Das ist eben das Schlimme: Die Person, die das aufgesetzt hat, hat nicht sehr sorgfältig über die Daten nachgedacht, aus denen der Algorithmus lernt. Dazu kommt: Manche Algorithmen erzeugen Entscheidungsregeln, die wir gut nachvollziehen können. Und andere erzeugen Entscheidungssysteme, die undurchsichtig sind. Wir nennen sie subsymbolisch, denn sie können Ihnen nicht in Symbolen und abstrakt erklären, wie sie zu ihren Entscheidungen gekommen sind, und man hat keine Chance, da hineinzugucken.

Wie bitte? Entwickler verstehen ihre eigenen Algorithmen nicht? Das klingt jetzt schon etwas nach Frankenstein...

Nein, den Algorithmus verstehen sie schon. Aber dieser baut ein System von Gleichungen auf und sagt: Ganz am Anfang kommen alle Bildpixel, das sind vielleicht 1024×1024 , hier rein. Das heißt, das sind Formeln mit 1024×1024 Variablen. Dann berechnen die etwas, und die Ergebnisse gehen in die nächste Schicht, in eine neue Riege von Formeln hinein. Und dann wieder in die nächste Schicht. Und da weiß kein Mensch mehr, warum was entschieden wird. Der Algorithmus, der das macht, der ist leicht verständlich, aber die Entscheidungen, die nachher aus den Daten herausgesogen werden, stecken in diesem riesengroßen Gleichungssystem. Das ist ein bisschen so, als wenn Sie mir erzählen, Sie mögen Kommissar Wallander. Dann habe ich keine Chance zu wissen, warum. Ich kann mir Teile der Begründung herleiten, wenn ich Sie besser kennenlerne – dann merke ich, Sie finden Schweden schön und mögen generell Kri-

mis, aber genau werde ich nie dahinterkommen. Sie selbst ja auch nicht. Und so ist es auch mit neuronalen Netzen. Man kann nur ganz lange zugucken, und dann weiß man irgendwann, oh, der denkt: Die Krauwatte bedeutet »nicht kriminell« und das T-Shirt ist für das System gleich »kriminell«.

Netflix kategorisiert mich auch und stellt Prognosen auf, welche Serien mir gefallen könnten – gehört das auch zu den bedenklichen Algorithmen?

Nein, eigentlich nicht. Wir haben für diese Frage ein Schadenstiefe-Modell erarbeitet, und ein Aspekt dabei ist: Gibt es andere Anbieter? Wenn es sie gibt und ich frei wählen kann, ist es tendenziell erst mal nicht so schlimm. Dabei müssen auch Quasi-Monopole mit bedacht werden. Denn wenn mein Arbeitgeber eine Software nutzt, um die Leistung der Arbeitnehmer zu bewerten, kann man natürlich sagen: Es gibt ja noch andere Arbeitgeber. Aber ich bin eben in dieser Firma, und das ist ein Quasi-Monopol, weil ich ohne erheblichen Aufwand da nicht rauskomme. Der zweite Punkt ist: Wie stark beschneidet mich ein Fehlurteil? Verschende ich nur eine halbe Stunde, weil mir die empfohlene Serie nicht gefällt, oder gehe ich zwei Jahre ins Gefängnis? Und dann gibt es noch Effekte, die auf der gesellschaftlichen Ebene erst sichtbar werden. Zum Beispiel wenn ein Suchmaschinen-Algorithmus oder ein soziales Medium mir Nachrichten serviert, die ich gar nicht sehen will, ist mein individueller Schaden nicht groß. Aber wenn das manipulativ eingesetzt würde, macht es einen Riesenunterschied, ob 80 Millionen eine Nachricht bekommen, die sie eigentlich nicht sehen wollten, die sie aber doch beeinflusst. Deswegen gibt es diese drei Dinge: Schaden auf individueller Ebene, Schaden auf gesellschaftlicher Ebene, der größer sein kann als der Einzelschaden, und die Frage, ob ich dem als Monopol unterworfen bin oder nicht.

Von diesen Empfehlungsalgorithmen fühlt man sich ja trotzdem sehr durchschaut. Beeinträchtigen die unsere Freiheit?

Das ist eine schwierige Frage. Die Frage ist ja, hatten wir jemals Entscheidungsfreiheit oder nicht? Ich antworte mal anders darauf. In der Community ist es kein großes Geheimnis, dass Spieledesigner und Webseitendesigner sehr gut zugehört haben bei den Psychologen, wenn es darum geht, worauf Menschen reagieren. Wer zum Beispiel ein Handygame spielt, nimmt unbemerkt an Tests teil, was am besten funktioniert, um die Leute länger am Spielen zu halten, und das wird ja angewendet und funktioniert. Insofern denke ich, dass wir beeinflussbar sind, das haben die psychologischen Studien der letzten Jahrzehnte deutlich gezeigt. Ich will hoffen, dass wir noch entscheidungsfrei sind, aber es scheint ziemlich viele Nebenbedingungen zu geben. Unter Stress sind wir weniger entscheidungsfrei als im entspannten Zustand. Wenn wir den ganzen Tag schon entschieden haben, sind wir nicht so frei in unserer Entscheidung, wie wir es sind, wenn wir gerade frisch aufgestanden sind. Unter dieser Vorbedingung setzen uns Spiele und Algorithmen noch mehr unter Druck, sie geben uns noch weniger Freiraum. Auch dadurch, dass sie uns Gruppen zuordnen. Das nenne ich ein algorithmisch legitimes Vorurteil. Denn es wird ja gesagt: Die Frau Zweig ist genau wie alle anderen 42-jährigen Mütter mit zwei Kindern, und deswegen kriegt die jetzt auch Pampers angeboten.

Was ist so schlimm daran?

Na ja, der Algorithmus weist mir eine recht bunt zusammengewürfelte Truppe zu, und ich habe keinerlei Einspruch, um zu sagen: Aber ich habe doch mit denen gar nichts gemein. Das ist beim Pampers-Beispiel egal, und nicht egal, wenn ich einer Gruppe von Kriminellen zugeordnet werde, die zu 80 Prozent rückfällig wird – ich aber zu den 20 Prozent gehöre, die nicht rückfällig werden. Wir verdammen das sowieso schon, aufgrund einer Gruppenzugehörigkeit verurteilt zu werden. Ich finde es aber noch schlimmer, wenn es eine Gruppe ist, mit der ich tatsächlich nichts gemeinsam habe, die mir nicht mitgeteilt wird, die ich nicht kenne, von der ich mich nicht abgrenzen kann.

Das würde heißen, dass Menschen doch nicht so berechenbar sind, wie es Algorithmen nahelegen?

Auch das ist eine wichtige und interessante Frage, auf die es zwei Antworten gibt: Im Großen und Ganzen sind wir viel berechenbarer, als wir denken. Das liegt daran, dass wir aus Einzeleigenschaften bestehen. Also: Ich lese gern Krimis, und das habe ich mit all den Leuten gemein, die sie auch gern lesen. Außerdem lese ich noch gern Programmiersprachenbücher und Liebesromane. In dieser Mischung bin ich individuell. Aber der Algorithmus interessiert sich normalerweise nicht für die Mischung, sondern versucht zu erraten, was das nächste Buch für mich sein könnte. Und dann sieht er, aha, sie liest Krimis, also bekommt sie den nächsten Krimi. Sie gehört zu den Programmiersprachenleuten, also bekommt sie auch das neue Programmiersprachenbuch. Und in diesem Sinne sind wir fantastisch vorhersehbar – gerade was Produkte angeht. Außerdem sind wir sehr gut vorhersehbar in unserem zeitlichen Verhalten, weil die meisten von uns sehr periodische Lebensläufe haben.

Und die zweite Antwort?

Wir sind vielleicht in unseren Einzelteilen und in unserem Einzelverhalten extrem vorhersagbar, aber interessant wird es dann natürlich an den Stellen, wo wir es nicht sind. Und davon gibt es immer noch genügend. Also ja, wir sind berechenbarer und weniger individuell, als wir es gern von uns denken. Aber wir sind nicht gut genug berechenbar in komplexen Fragen wie der Rückfälligkeitsquote bei Kriminellen, Arbeitsleistungsfähigkeit, vielleicht auch Kreditwürdigkeit, das kann man nicht algorithmisch machen.

Gibt es noch etwas, das Algorithmen nicht können?

Der Algorithmus wird mir nie das eine Schmuckstück auswählen können, das ich besonders schön finde. Aber unter Produkten wie Büchern eines auszusuchen, was ich mag, das ist ganz einfach. Aber da gibt es einen Creepyness-Faktor, wir finden es etwas gruselig, dass es eine Maschine war, die das für mich

ausgesucht hat. Das ist auch der Grund, warum wir computergenerierte Kunst niemals akzeptieren werden außer als Kuriosität. Weil dahinter eben nicht die Lebenszeit eines Menschen steht – deswegen war es bei Andy Warhol schon kritisch, als er seine Serien hat drucken lassen.

Was sind in Ihren Augen Bereiche, in denen der Einsatz von Algorithmen sinnvoll ist?

Was ich gerade besonders spannend finde, ist das autonome Fahren. Aber im Moment kämpfen wir damit, dass man die Algorithmen noch leicht austricksen kann. Aus einem Stoppzeichen kann man zum Beispiel mit wenigen Aufklebern, die wir als Menschen einfach ausblenden, für die Maschine ein 80-km/h-Zeichen machen. Das ist natürlich nicht günstig. Aber auf autonomes Fahren hoffe ich. Bei 3000 Toten und 100 000 Verletzten im Jahr, weil Menschen müde fahren, mit Handy am Ohr oder einfach Fehler machen, könnte es viele Menschenleben retten. Aber auch da werden wir es wahrscheinlich viel gruseliger finden, wenn dann ein solches Auto jemanden überfährt. Wir werden das der Maschine ganz anders anlasten als einem Menschen.

Über den »Todesalgorithmus«, der in entsprechenden Situationen entscheiden muss, wer überfahren werden soll, wird ja viel diskutiert. Was denken Sie darüber?

Das Auto muss für jede Situation eine Regel haben, was es tut. Und natürlich werden die Regeln grundsätzlich lauten: Weiche einem Hindernis aus, wenn es möglich ist. Und dann gibt es Situationen, in denen die Frage auftaucht: Welchem Hindernis ausweichen, wenn es ein Baum, drei Kinder, zwei Omas, drei Dackel, was auch immer, sind?

Wie beim Trolley-Problem aus der Philosophie. Genau.

Sie wirken etwas genervt davon.

Ja, furchtbar! Weil es nicht die Situation ist, die diese

Autos milliardenfach entscheiden werden müssen. Es werden sehr wenige Fälle sein, und es sind Fälle, die wir im Moment im Gesetz auch nicht regeln. Wenn einer von uns in eine solche Situation käme, steht in keiner Straßenverkehrsordnung, was wir tun sollen. Sondern ein Gericht entscheidet im Nachhinein, ob wir fahrlässig waren oder nicht. Und das Gute an den Maschinen ist, sobald wir entschieden haben, das war jetzt falsch, können wir sie sofort umprogrammieren. Das ist bei uns Menschen nicht so. Wenn ich dafür bestraft werde, heißt das nicht, dass Sie das in Zukunft nicht tun. Diese Diskussion bekommt viel zu viel Gewicht.

In Ordnung, aber wer haftet denn, beziehungsweise wer trägt, nun wieder philosophisch gefragt, die Verantwortung? Wenn man einen Algorithmus programmiert, ist man dann für ihn verantwortlich?

Wenn ich einen Algorithmus programmiere, der fest in einem Navigationsgerät verbaut ist und der nur dafür gebraucht wird, dann habe ich volle Kontrolle über alles. Über die Software und den Kontext. Da kann man Autorenschaft und Akteursschaft noch sehr gut zuordnen. Doch auch da sind wir im Gespräch mit Philosophen, ob das wirklich so ist. Aber ich habe zum Beispiel einen Algorithmus entwickelt, den ich allein schon für drei sehr verschiedene Fragestellungen eingesetzt habe: Es ist ein Empfehlungsalgorithmus für Filme, dann stellten wir fest, dass er sich aber auch eignete, um eine 30 Jahre alte Theorie der Meerforschung zu widerlegen. Und wir haben damit eine Brustkrebstherapie identifiziert. Die weiteren Einsatzmöglichkeiten sind nur durch die Fantasie meiner Mitforschenden begrenzt. Wenn das der Fall ist, wenn ich als Designerin nicht weiß, in welchem Kontext mein Algorithmus benutzt wird, dann würde ich sagen, ist man auch nicht dafür verantwortlich. Interessant wird auch das wieder bei den lernenden Algorithmen, da gab es mal diesen Chatbot von Microsoft. Der hieß Tay, war auf Twitter aktiv und sollte lernen, sich mit Menschen zu unterhalten. Ein

paar Idioten haben sich verabredet, ihn anzustacheln, und ihm rassistische und sexistische Tweets geschickt – und siehe da, der Bot lernt es und kann auch ein Rassist und ein Sexist sein. Dann war er innerhalb von 24 Stunden vom Netz. Da hätte ich persönlich mir etwas mehr Standhaftigkeit gewünscht von den Designern, denn ich finde, das Ding hat das getan, was es tun soll. Es sollte lernen. Ich hätte mir eher gewünscht, dass man sagt: So liebes Internet, jetzt zeigt doch bitte mal, dass ihr das auch wieder zurückziehen könnt. Darüber diskutieren wir gerade viel. Die lernenden Algorithmen, die in der Interaktion mit Menschen dumme Dinge tun.

Eines der zumindest unerwünschten Dinge, die Algorithmen in Interaktion mit Menschen tun, ist zu lernen, was die Menschen mögen, und ihnen das immer wieder zu empfehlen und so die Filterblasen aufzubauen. Was sagen Sie dazu?

Das ist auch so eine Pseudo-Diskussion.

Na ja, aber man könnte doch grundlegend infrage stellen, ob es gut ist, dass man Dinge vorgeschlagen bekommt, die so sind, wie das, was man schon kennt.

Aber wer lebt denn ganz allein im Wald mit einem 1-Gigabit-Anschluss? So einen Menschen gibt es doch gar nicht! Wir sind ja noch vielen anderen Dingen ausgesetzt, und ich denke, dass wir dafür verantwortlich sind, uns auch noch im *real life* zu treffen und uns auszutauschen. Also dieser Wert geht nicht kaputt, nur weil es diese Systeme gibt.

Wie schätzen Sie das in Bezug auf Facebook ein? Es wurde ja vielfach kritisiert, dass der Algorithmus die Filterblasen verstärkt und dadurch zusammen mit Fake News etc. die Trump-Wahl gefördert hätte, sodass es schon Forderungen gab, Facebook sollte seinen Algorithmus offenlegen.

Wer das fordert, irrt sich in zwei Punkten. Erstens, wenn ich nicht will, dass andere meinen Code verstehen, kann ich den beliebig kompliziert machen, dass

kein Mensch mehr versteht, was hier los ist. Oder dass er zumindest sehr viele Personenstunden braucht, um das herauszufinden. Zweites, wenn es derart viele Personenstunden sind – die Pornoindustrie und andere haben immer das Geld dafür. Das heißt, die Wahrscheinlichkeit, dass die das verstehen, bevor wir das verstehen, ist riesengroß. Und dann hat man ein System völlig kaputt gemacht. Die fangen an, mit dem System zu spielen, bis sie ihre Websites ganz oben haben. Die Phishing-Seiten, die Pornoseiten, Seiten, die man nicht sehen will. Deswegen ist Offenlegung des Codes sinnlos. Offenlegung gegenüber Experten könnte sinnvoll sein. Da komme ich wieder zurück zu meinem Schadenstiefe-Instrument. Ich persönlich finde, wenn es ein staatliches Monopol gibt, wie etwa Geheimdienst-Algorithmen, die Individuen und der Gesellschaft erheblich schaden können, dann sollten Wissenschaftler in die Daten oder das Entscheidungssystem hineingucken können.

Aber nach Ihrem Schadenstiefe-Modell müsste der Facebook-Algorithmus doch durchfallen. Es ist ein Quasi-Monopol, und der Schaden für die Gesellschaft könnte groß sein.

Ja, so wie wir es nutzen, ist es ein Quasi-Monopol, und wenn es um Politisches geht, könnte der gesellschaftliche Schaden enorm sein. Man könnte also sagen: grundsätzlich keine politische Werbung auf Facebook. Aber wer entscheidet, was politische Werbung ist? Und wenn man sagt, persönliche Meinung darf bestehen bleiben, kommen die Bots. Da muss jemand entscheiden, wer Bot ist und wer nicht. Oder wir sagen: soziale Medien ab heute nur noch mit Kuschelthemen. Du darfst dich über alles unterhalten, nur nicht über Politik. Dann gäbe es keine freie Meinungsäußerung mehr in dem Netzwerk. Aber warum sind soziale Medien so gefährlich geworden, etwa in diesem ganzen Trump-Fake-News-Zirkel? Ich denke, das liegt auch daran, dass wir alle zu One-to-many-Distributoren geworden sind. Früher konnte ich meine Verschwörungstheorien mit meinem näheren Umfeld austauschen, heute kann ich innerhalb von

Biografie

Katharina Zweig, 42, ist Professorin und leitet das »Algorithm Accountability Lab« an der TU Kaiserslautern, wo sie 2012 den deutschlandweit einzigartigen Studiengang »Sozioinformatik« mit aufgebaut hat. Er beschäftigt sich damit, was passiert, wenn Menschen über Software miteinander interagieren, also den Wechselwirkungen von Gesellschaft und Informatik. »Wir müssen Mensch und Software als komplexes System zusammendenken«, sagt Zweig, »es nützt nichts, wenn wir nur die Software angucken oder nur das soziale System. Durch die Interaktion passieren Dinge, die man sich vorher nicht vorstellen konnte.« Ein einfaches Beispiel ist die Google-Suchervollständigung, die aus Eingaben von Menschen lernt und dann auf sie zurückwirkt.

Mehr Informationen unter:

aalab.cs.uni-kl.de

Mehr über den Studiengang unter:

www.cs.uni-kl.de/studium/studiengaenge/bm-si/sp.ba

Sekunden an Tausende Leute herankommen. Deshalb schlage ich vor, man könnte zum Beispiel die Retweet-Geschwindigkeit von Nachrichten runtersetzen. Damit, falls die Information falsch ist, sich die betroffene Person dagegen wehren kann oder die falsche Information sich nicht rasend schnell verbreitet. Man könnte an dieser Stelle aus dem Internet die Geschwindigkeit etwas herausnehmen, das könnte einen Teil der Probleme lösen. Wir müssen kreativ werden abseits von klassischen Verboten.

Können Algorithmen menschliche Fehler auch bei so etwas wie Diskriminierung ausbügeln und vielleicht fair und objektiv aus den Bewerbern für einen Job auswählen?

Fair und objektiv zu sein, ist auch für die Maschine schwierig, weil auch sie aus der Vergangenheit lernt. Was ist, wenn die letzte Personalmanagerin einfach keine Rothaarigen mochte? Findet man darauf Hinweise in den Daten, lernt die Maschine das mit. Das heißt, man müsste sich ganz sicher sein, dass die Da-

ten, die man vorn reingibt, diskriminierungsfrei sind. Und dafür möchte ich bei niemandem die Hand ins Feuer legen. Und es kommt erschwerend hinzu, dass die Maschine wahrscheinlich lernt für alle Firmen dieser Welt: 28-jährige BWLer mit einem halben Jahr Auslandserfahrung sind die Besten. Weil das einer der häufigsten Studiengänge ist, weil das grundsätzlich oft smarte Leute sind, wenn die ihr Studium schnell gemacht haben und noch Zeit hatten fürs Ausland. Dann lässt das auf vieles hindeuten, einen guten elterlichen Hintergrund und so weiter. Und dann kriegen wir alle nur noch Mainstream. Denn was die Maschinen nicht gut können, wozu sie auch gar nicht da sind, ist, den Exoten zu finden, der trotz eines Lebenslaufs, der überhaupt nicht zur Firma passt, eine totale Bereicherung wäre. Denn im Wesentlichen lernen Maschinen das, was häufig ist. Einen Querdenker findet die Maschine nicht.

Also doch wieder die Menschen?

Ich glaube, es ist okay, wenn Algorithmen die 30 Prozent schlechtesten Bewerbungen raussuchen. Da kann man nicht viel falsch machen. Aber ich wäre deutlich vorsichtiger, wenn eine Arbeitgeberin damit die besten zehn Prozent raussuchen wollen würde. Es gibt wahrscheinlich keinen guten Grund, Personalentscheidungen einem Algorithmus zu überlassen.

Das erinnert mich daran, dass Sie in einem Interview mal gesagt haben: »Computer können uns nicht die Entscheidung abnehmen, wie wir leben wollen.«

Das stimmt. Allein schon, was diesen kleinen Bereich der Personalfragen angeht, steckt dahinter ja die philosophische Frage: Was ist überhaupt eine gute Entscheidung? Wollen wir Leute finden, die möglichst viele Jahre da bleiben? Welche, die möglichst schnell aufsteigen? Oder welche, die Diversität abdecken? Wer soll das denn entscheiden? Die Maschine kann es auf jeden Fall nicht. Spätestens an dieser Stelle braucht es Menschen, die wissen, dass sie hier ethische Entscheidung treffen. ■

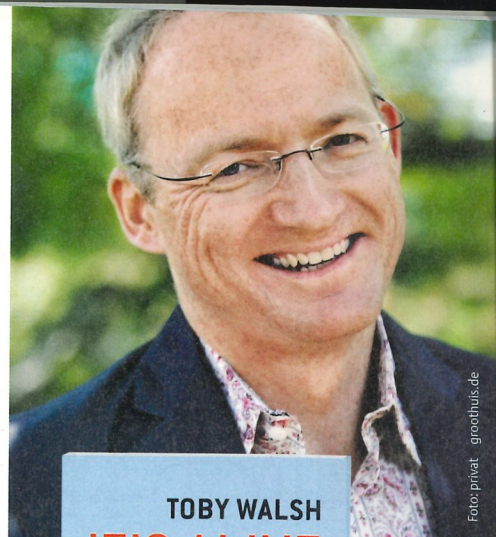
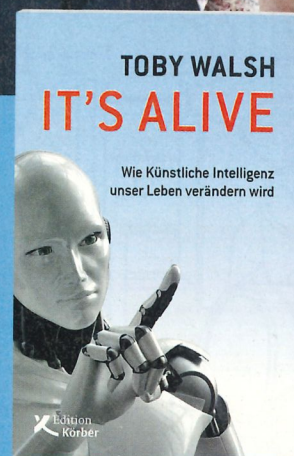


Foto: privat - groothuis.de



ISBN 978-3-89684-266-4

352 Seiten | Klappenbroschur

€ 18,- (D) | € 18,50 (A) | CHF 25,90

»Wenn Programme sich unserem Niveau von Intelligenz annähern, neigen wir dazu, sie für intelligenter zu halten, als sie tatsächlich sind.«

Der australische KI-Experte Toby Walsh erzählt anschaulich und mit trockenem Humor, wie die superschlauen Maschinen schon heute die Gesellschaft, Wirtschaft, ja, sogar uns selbst verändern.

Die Zukunft hat begonnen!

**Edition
Körber**

www.edition-koerber.de