



2021/05 Networld

<https://shop.jungle.world/artikel/2021/05/automatisierte-diskriminierung>

In Stanford hat ein Algorithmus zur Priorisierung der Impfungen versagt

Automatisierte Diskriminierung

Von **Enno Park**

Im kalifornischen Stanford ist ein Algorithmus zur Verteilung von Impfungen kläglich gescheitert. Doch was manchen als neuerlicher Beweis dient, dass Algorithmen nicht zu trauen sei, zeigt letztlich nur, dass deren Fehler menschengemacht sind.

Die Universitätsklinik Stanford im US-Bundesstaat Kalifornien hat so viele Mitarbeiter wie manch eine Kleinstadt Einwohner. Einschließlich Kinderklinik, Forschungsabteilungen und Lehrbetrieb arbeiten dort rund 30 000 Menschen. Zu entscheiden, in welcher Reihenfolge sie gegen Sars-CoV-2 geimpft werden sollen, ist schwierig, wenn, wie Mitte Dezember, zunächst nur 5 000 Impfdosen zur Verfügung stehen.

Also entwickelte eine Arbeitsgruppe einen Algorithmus, der festlegte, in welcher Reihenfolge die Beschäftigten geimpft werden sollten. Das Ergebnis zog wütende Proteste von Assistenzärztinnen und -ärzten nach sich. Sie gehören zu denen, die an der Universitätsklinik am häufigsten in direkten Kontakt mit Patientinnen und Patienten kommen, viele auch mit solchen, die an Covid-19 erkrankt sind. Dennoch sollten laut Algorithmus Verwaltungsangestellte im Homeoffice und hochrangige Ärzte mit wenig Patientenkontakt vor ihnen geimpft werden. Von den rund 1 300 Assistenzärzten, die in der Klinik arbeiten, sollten nur sieben an der ersten Impfrunde teilnehmen.

Von den rund 1 300 Assistenzärzten, die in der Universitätsklinik Stanford arbeiten, sollten nur sieben an der ersten Impfrunde teilnehmen.

Der Fall gilt als neuestes Beispiel in einer Reihe von Fehlentscheidungen, die von Computer-Algorithmen getroffen wurden. An ihm lässt sich sehr gut nachvollziehen, wie es zu solchen Fehlern kommt. Die Expertinnen und Experten, die den Algorithmus entwickelten, berücksichtigten neben dem Einsatzort der Beschäftigten gängige Risikofaktoren, die die Wahrscheinlichkeit erhöhen, dass ein Mensch schwer an Covid-19 erkrankt oder stirbt, im wesentlichen das Alter der Beschäftigten. Das scheint medizinisch und ethisch geboten. Es führte aber dazu, dass die häufig eher jungen Assistenzärzte benachteiligt wurden gegenüber meist älteren Verwaltungsangestellten sowie

höherrangigen Ärzten – wegen ihres Alters und offenbar auch, weil ihnen kein fester Einsatzort zugewiesen worden war. Die vorrangige Orientierung am Alter ergab in diesem Fall wenig Sinn, lassen sich doch die Verwaltungsangestellten ins Homeoffice schicken, während Assistenzärzte nicht nur gefährdeter sind, sich anzustecken, sondern das Virus auch von Patient zu Patient tragen können.

Solche Fehler führen dazu, dass Algorithmen Menschen diskriminieren, obwohl sie nicht absichtlich so programmiert wurden. Ein weiteres Beispiel aus dem US-amerikanischen Gesundheitssystem: Wie ein Algorithmus dort schwarze Patienten systematisch benachteiligte, beschrieb Ziad Obermeyer, Professor für Gesundheitspolitik und -management an der Universität Berkeley in Kalifornien, in einer 2019 in der Fachzeitschrift *Science* veröffentlichten Studie. Bei gleich schwerer Erkrankung erhielten Schwarze weniger oft eine Zusatzbehandlung für Risikopatienten als Weiße. Die Analyse ergab, dass das Problem auf einem Denkfehler beruhte: Die Programmierer des Algorithmus hatten angenommen, dass Patienten, die hohe Kosten verursachen, eher Risikopatienten sind und eine zusätzliche Behandlung benötigen. Da schwarze Arbeitnehmer in den USA häufiger schlecht bezahlte Arbeitsplätze haben als weiße und deshalb schlechter krankenversichert sind, können sie weniger Geld für medizinische Behandlungen ausgeben. Sie verursachen weniger Kosten als weiße Patienten und wurden deshalb vom Algorithmus seltener als Risikopatienten ausgewiesen.

Aber nicht nur fehlerhafte oder ungeeignete Algorithmen verursachen Diskriminierung, sondern auch mangelhafte Daten. Die Meriten, die sich Forscher und Forscherinnen erwerben können, werden unter anderen mit dem so genannten H-Index gemessen. Der von dem Physiker Jorge E. Hirsch entwickelte Index bewertet Wissenschaftler anhand der Anzahl von Veröffentlichungen und Zitationen. Er hat starken Einfluss auf die Karrierechancen vieler Wissenschaftler. Allerdings ergibt sich ein Problem aus Namensänderungen, wie sie insbesondere nach Heiraten vorkommen. Ältere Publikationen, die unter dem Geburtsnamen veröffentlicht wurden, werden dann häufig nicht korrekt in den H-Index eingerechnet, so dass dieser zu gering ausfällt. Da bei Heiraten überdurchschnittlich häufig Frauen ihren Nachnamen ändern, produziert der Algorithmus sexistische Diskriminierung.

Oftmals werden solche Fehler auch mit sogenannter künstlicher Intelligenz (KI) in Verbindung gebracht. Solche selbstlernenden Systeme, die beispielsweise Gesichter erkennen, gesprochene Sprache übersetzen oder das Wetter vorhersagen können, geben problematische Ergebnisse aus, wenn ihre Mustererkennung mit problematischen Daten trainiert wird. Das lernte Microsoft, als das Unternehmen 2016 einen Twitter-Bot namens Tay ins Leben riefen. Dieser sollte selbständig anhand der Alltagskommunikation hinzulernen. Im Handumdrehen fanden sich Trolle, Nazis und Nazitrolle, die Tay mit einschlägigen Inhalten fütterten. Innerhalb kurzer Zeit gab Tay sexistische, rassistische und antisemitische Äußerungen von sich und wurde nach nur 16 Stunden abgeschaltet.

Meldungen über selbstlernende Computersysteme, die Fehler machen, gibt es regelmäßig. Sie führen zu dem Eindruck, dass KI an sich ein Problem sei. Genauso heißt es oft, wenn eine Panne passiert, der Algorithmus sei schuld. Doch damit werden künstliche Intelligenz und Algorithmen zu Handelnden erhoben und manchmal zu diffusen, von normalen

Menschen kaum zu durchschauenden Wesenheiten erklärt. Wann immer man sagt, ein Algorithmus habe dies getan oder eine künstliche Intelligenz jenes, verschleiert man, dass Menschen diese Systeme programmiert und trainiert haben.

KI ist nichts weiter als ein statistisches Verfahren zum Errechnen von Wahrscheinlichkeiten, also ein komplexer Algorithmus. Und ein Algorithmus wiederum ist nichts weiter als eine Rechenanweisung. Wenn die Süddeutsche Zeitung in einem Artikel schreibt: »In den USA entscheidet ein Algorithmus, wer zuerst an den Corona-Impfstoff kommen darf«, dann ist das zwar nicht falsch, aber in Deutschland entscheidet eben auch ein Algorithmus über die Impfreihefolge, nämlich die Regeln, die das Bundesgesundheitsministerium aufgestellt hat. Ob ein Computersystem oder ein Behördenapparat diesen Regeln folgt, ist zunächst einmal egal, und wer Ersteres meint, sollte sicherheitshalber immer von »Computer-Algorithmen« sprechen. Nüchtern betrachtet handelt es sich bei KI und Algorithmen nur um automatisierte Entscheidungssysteme.

Solche Systeme sind nicht deswegen problematisch, weil sie nach Algorithmen arbeiten und dafür gar KI benutzt wird, sondern weil sie so schwer zu kontrollieren sind. Menschen neigen zudem dazu, Angaben, die ein Computer macht, einfach hinzunehmen. Aus einer Reihe von Gründen ist es sehr schwer, die Sachbearbeiter einer Bank, Behörde oder Firma davon zu überzeugen, dass falsch oder ungerecht ist, was da auf ihrem Computerbildschirm steht. Sich einfach danach zu richten, erfordert von ihnen weniger Energie, als auf Fehlersuche zu gehen, dabei an Produktivität einzubüßen und sich - womöglich beim Vorgesetzten unbeliebt zu machen. Der normale Handyverkäufer hinterfragt die schlechte Schufa-Bewertung eines Kunden nicht, sondern verweigert kurzerhand den Vertrag.

Um dagegen vorzugehen, braucht es keine »Algorithmenethik«, wie häufig gefordert, sondern Einspruchs- und Kontrollmöglichkeiten. Dass automatische Entscheidungssysteme viele Probleme erzeugen, liegt auch daran, dass Entscheidungen automatisiert werden, die man nicht so einfach treffen kann. Die Debatten darüber, ob selbstfahrende Autos im Ernstfall den Rentner oder die Frau mit dem Kinderwagen überfahren sollen, enden nicht, weil es auf die Frage keine allgemeingültige Antwort gibt. Die Versuche, solche Entscheidungen in Algorithmen zu fassen, machen das oftmals erst deutlich. Ein selbstlenkendes Auto muss nun einmal irgendwie programmiert werden, aber wenn es dann per Algorithmus Entscheidungen trifft, ist das auch nicht recht.