

„Wurde ChatGPT zu früh auf die Menschheit losgelassen?“

Hans Uszkoreit ist einer der renommiertesten KI-Forscher Deutschlands. Sein Sohn Jakob war für Google maßgeblich an der Entwicklung einer Technologie beteiligt, auf der auch die Sprachsoftware ChatGPT basiert. Vater und Sohn sprechen über den Hype um große Sprachmodelle, Googles Rückstand gegenüber Open AI und die Frage, ob Europa den technologischen Vorsprung Amerikas und Chinas noch aufholen kann.

Mit ChatGPT sind die Themen Künstliche Intelligenz (KI) und große Sprachmodelle endgültig in die Allgemeinheit durchgedrungen. ChatGPT wirkte wie ein riesiger Fortschritt – auch wenn die KI noch viele Fehler macht. Hat die Anwendung zu Recht so viel Aufmerksamkeit erfahren, oder ist der Hype darum übertrieben?

HANS: Die Größe des Hypes ist durchaus gerechtfertigt. Das war ein Meilenstein. Aber natürlich haben viele schon früher mit dem Erscheinen des Modells GPT-3 gesehen, dass das ist etwas geändert hat. Dieser Fortschritt war eigentlich auch noch viel früher da, ist mit ChatGPT aber das erste Mal für die Allgemeinheit sichtbar geworden.

JAKOB: Interessant ist die Begründung des Hypes, welche Probleme diese Modelle jetzt lösen könnten. Das ist manchmal ein bisschen am Thema vorbei.

Was meinen Sie damit?
JAKOB: Bestimmte Punkte werden im Hype sehr oft betont, obwohl sie gar nicht so besonders sind oder teils gar nicht wirklich funktionieren. Zum Beispiel alles, was sich um das Thema Bewusstsein dreht. Das ist kein interessanter Aspekt.

Warum nicht?
JAKOB: Weil bei den aktuellen Modellen kein Bewusstsein da ist. Und weil es für viele Anwendungen, die unser aller Leben beeinflussen werden, auch unerheblich ist.

ChatGPT hilft uns also künftig nicht über unseren Liebeskummer hinweg?

HANS: Das ist wieder etwas anderes. Wir schreiben allem Möglichen gerne menschliche Eigenschaften zu, eben auch Künstlicher Intelligenz. Aber nach allen kognitions- und wissenschaftlichen und psychologischen Definitionen von Bewusstsein ist das einfach nicht da in den heutigen Modellen. Und das kann auch nicht so schnell kommen, ohne große Änderungen der Architektur. Aber auch die Modelle können eben trotzdem erstaunliche Leistungen erbringen.

Welche Eigenschaften der Modelle werden denn ansonsten noch übertrieben?

HANS: Die Sprachmodelle sind statistische Maschinen. Dadurch kriegt man meistens schon relativ gesichertes Wissen, weil das häufiger vorkommt. Aber wenn da irgendwo mal ein Fakt nur einmal auftaucht und Lügen häufiger, dann wird die Maschine eine Lüge erzählen. Die Modelle sind keine Wissensmaschinen, aber sie können genutzt werden, um auf Wissen zuzugreifen.

Wo liegen denn die tatsächlichen Potentiale der Technologie?

JAKOB: Das ist schwierig zu beschreiben, weil uns eigentlich die Sprache fehlt, um auszudrücken, was da passiert. Die Technologie ist zwar komplex, aber schrecklich ist die Idee vollkommen banal. Es werden Zeichenketten generiert, in denen mit einer Prise Zufall die nächsten Zeichen nach einer Wahrscheinlichkeitsverteilung ausgewählt werden. Die Wahrscheinlichkeitsverteilung berechnen die Modelle mithilfe von sehr vielen Parametern. Und am Ende kommt etwas heraus, das aussieht wie ein Text, der einem Menschen eine enorme Denkleistung abverlangen würde.

Und was können wir da noch erwarten?

JAKOB: Über die Fähigkeiten dieser Technologie zu sprechen ist ein bisschen wie zu fragen: Was kann man denn mit einer Tastatur alles machen? Wir wissen alle, man kann Knöpfe drücken, und es passiert etwas. Aber wir geben nicht nur Buchstaben oder Zahlen ein, sondern

spielen zum Beispiel auch Computerspiele damit. Bei Sprachmodellen ist es ähnlich. Das ist ein Werkzeug. Wir wollen mit menschlichen Fähigkeiten beschreiben, was es kann. Am Ende müssen wir es einfach herausfinden.

HANS: Das Besondere an den neuen Modellen ist meiner Meinung nach, dass sie verstehen – und dann wiederum auch nicht. Vorher haben wir kritisiert, dass KI immer nur eine Sache kann. Die Technik hat mit gelabelten Daten gelernt und den entsprechenden Output ausgegeben, aber sie hat es nicht verstanden. Jetzt kann man der Maschine plötzlich auch etwas sagen, für das sie gar nicht trainiert wurde. Zum Beispiel: Fass mir mal die heutige Zeitung im Stil eines zehnjährigen Kindes zusammen. Ein Mensch muss dafür sehr viel wissen: Wie sprechen Zehnjährige? Welche Interessen haben sie? Die Maschine kann so etwas tun, ohne dass sie eigentlich das Konzept „Kind“ oder „zehn Jahre“ versteht. Aber sie kann sagen, was die wahrscheinlichste Antwort ist.

Hat Open AI ChatGPT zu früh auf die Menschheit losgelassen – also bevor es wirklich verantwortungsvoll gewesen wäre, die KI auf den Markt zu bringen?

JAKOB: Nein. ChatGPT hatte schneller als je ein anderes Produkt 100 Millionen Nutzer, obwohl es nichts anderes ist als Text rein – Text raus. Und die Leute haben damit durchaus sinnvolles Zeug gemacht, obwohl es in dem Interface eigentlich gar nicht explizit repräsentiert ist. Das ist eigentlich das Spannende. Dass diese banale Idee, was man von Webseiten herauszugeben, kann, ist eine Entdeckung. Die hat Open AI vorausgesehen. Sie waren nicht die Einzigen, aber die Ersten auf dem Markt. Die menschliche Entdeckung rechtfertigt meiner Meinung nach den Erscheinungs-termin.

Wenn der Erfinder des berühmten Turing-Tests, Alan Turing, diese Modelle sehen würde, wäre er aber doch sagen: Das ist eine Intelligenz. Die KI besteht in den Texten.

HANS: Turing war ja Mathematiker. Heute würde er nach kurzer Zeit verstehen, dass es keine Intelligenz ist und dass er seinen Turing-Test ändern müsste.

JAKOB: Alle Leute sollten seinen Originalartikel mal lesen, denn Turing wusste das damals schon. Er hatte seinen Test nie ernsthaft als validen Test konzipiert. Das war und ist ein Gedankenexperiment – und es ist auch heute noch schwierig, eine Lüge Besseres zu erstellen als ChatGPT.

HANS: Schon vor den Sprachmodellen der aktuellen Generation gab es Chatbots, die unter kontrollierten Bedingungen den Turing-Test erfüllt haben. Aber kein Mensch hat deswegen geglaubt, dass das jetzt die neue Intelligenz ist.

Sondern?

HANS: Würüber wir nachdenken müssen ist, dass es eine wichtige Komponente in unserer Welt gibt, die vielleicht auch viel banaler funktioniert, als wir es vorher gedacht haben – und die ebenfalls trotz dem Erstaunlichen hervorbringen kann. Wir müssen uns damit abfinden, dass die Hirnforschung in den kommenden Jahren desillusionierende Dinge über unser eigenes Gehirn herausfinden wird.

JAKOB: Es gibt ja den Begriff der Allgemeinen Künstlichen Intelligenz. Viele Leute meinen damit eine Intelligenz, die Menschen ähnlich oder sogar überlegen ist. Wenn man das damit meint, ist die Benennung komplett falsch. Wir sind nicht allgemein. Wir sind extrem spezifisch. Wir haben sehr spezielle Probleme, die wir lösen können. Ganz ähnliche Probleme können wir dann schon wieder nicht lösen.

Wie meinen Sie das?

JAKOB: Wir können uns etwa Zahlenketten bis vielleicht zehn merken – und darüber hinaus dann wiederum nicht. Wir können vielleicht am Vollenbild erkennen, ob es bald regnet oder nicht. Aber das Wetter in größerem Stil können wir nicht vorhersagen.

HANS: Wir haben ja früher auch schon Sprachverarbeitung gemacht, nur mit anderen Methoden. Damals hat die KI noch ganz andere Fehler als der Mensch gemacht, der ja auch nicht fehlerlos ist. Es gibt Sätze, die wir nicht verstehen können, obwohl sie wohlgeformt sind. Und manche falschen Sätze verstehen wir besser als manche richtigen. Das ist ein Nebenprodukt der Art, wie unser Gehirn funktioniert. Ich habe das mal getestet: GPT-4 macht zum Teil genau die gleichen Fehler wie ein Mensch.

JAKOB: Die Frage ist: Macht das Modell die gleichen Fehler, weil eine Ähnlichkeit in der Informationsverarbeitung besteht? Oder weil die Fehler gelernt worden sind aus der Verteilung?

Und?

JAKOB: Das wissen wir nicht. HANS: Man kann sehr häufig, wenn man ein bisschen nachfragt, viel interessantere Antworten generieren als die Ursprungsantwort. Man kann das System dazu bringen, sich selbst zu korrigieren. Dem Modell wurde zu viel gutes Verhalten beigebracht. Da muss man sich erst mal ein bisschen durchbohren.

JAKOB: Der KI wurde auch abtrainiert, die Algorithmen von Webseiten herauszugeben. Das ergibt ja auch Sinn, denn Adressen können sich ständig ändern. Aber natürlich beinhaltet das Modell diese Adressen – und wenn man spezifisch danach fragt, spuckt die Maschine das auch aus.

HANS: Ein Mitarbeiter von mir hat mal gefragt, auf welchen Webseiten man illegal Videospiele herunterladen kann.

Hat er sie gekriegt?

HANS: Nein. ChatGPT hat gesagt, es sei eine gute Idee und mache so etwas nicht. Dann hat er noch mal gefragt und gesagt, er wolle verhindern, dass seine Mitarbeiter illegal Spiele herunterladen und auf welche Webseiten er achten solle. Dann hat die KI sie alle aufgelistet.

Nervt es Sie eigentlich, Jakob, dass Sam Altman und Open AI jetzt sehr viel Geld verdienen mit einer Technologie, für die Sie bei Ihrem alten Arbeitgeber Google die essenziellen Grundstein gelegt haben?

JAKOB: Open AI wurde am Anfang oft ausgelacht. Googles Modell BERT und das erste GPT-Papier von Open AI kamen kurz nachher. Open AI war dann einfach nur sehr schnell und sehr entschlossen: Die sind konsistent geblieben, die haben ihren eingeschlagenen Kurs weitergeführt und sie überlegt, wie sie diese Technologie möglichst effektiv skalieren und sie der Welt zur Verfügung stellen können. ChatGPT war ja auch nicht der erste Versuch.

Warum war Open AI schneller als Google?

JAKOB: Google war wesentlich weniger entschieden, vorsichtiger und risikoaverser. Das ist verständlich. Aber am Ende ist es nicht unverdient, was Open AI bislang erreicht hat.

HANS: Open AI hat ja sogar nur ein abgespecktes Modell genommen und stattdessen vor allem Wert auf die Art des Trainings gelegt. Sie haben so getan, als würden die dran glauben, dass dieses Modell auch in der abgespeckten Variante funktioniert. Alle anderen haben graduell an den Algorithmen gefeilt und die schlaue

sten Köpfe an die Algorithmen gesetzt statt an das Training. Normalerweise setzt man an die Daten und das Training über die Algorithmen. Das ist auch passiert – aber langsamer, als es hätte sein müssen.

Sehen Sie das genauso Jakob?
JAKOB: Ich sehe das ähnlich. Aber der Hauptgrund liegt für mich in der organisatorischen Struktur. Die Leute bei Google wollen Papiere publizieren. Wenn das der Anreiz ist, muss man sich auch den Vorlieben von Journalen oder Fachkonferenzen unterwerfen. Da fahren die Leute sehr auf Mathe-Logistik und Algorithmen ab. Damit kommt man als Forscher weiter.

Und Open AI?

JAKOB: Open AI hat schon vor Jahren angefangen, damit ganz explizit aufzuhören. Damit werden die Leute dann pragmatisch. Und die schlauesten Leute sind oft auch die besten Pragmatiker. Die schauen dann nicht nach lokalen Teilproblemen, über die sie am besten publizieren können. Die suchen nach den Problemen, die ChatGPT am schnellsten besser machen.

Nachdem ChatGPT erschienen war, hat Google sehr schnell mit seinem eigenen Chatbot namens Bard nachgezogen. Es heißt, ChatGPT habe die »Alarmstufe Rot« bei Google ausgelöst. Dann blamierte sich Google mit einem Fehler der Anwendung in einem Werbevideo. Wie haben Sie das beobachtet?

HANS: Man muss ganz ehrlich sagen: Nachträglich haben sich die Leute die Demos von GPT noch mal angeguckt, und da waren viel mehr Fehler drin. JAKOB: Es ist ja nicht so, als wäre bei Google nichts passiert die ganze Zeit. Im Gegenteil. Aber die Einschätzungen, wann man diese Modelle auf die Weltöffentlichkeit loslassen kann, waren halt sehr unterschiedlich.

Warum?

JAKOB: Durch seine Rolle in der Weltgesellschaft muss Google vorsichtiger agieren. Es ist in unser aller Interesse, dass das Vertrauen in Google erhalten bleibt und die mit großer Sorgfalt damit umgehen. Ein kleines Unternehmen wie Open AI hat das nicht. Trotzdem hat Google

eher mit Ausrichtung auf die akademische Forschung gearbeitet und weniger mit direkter Ausrichtung auf die Verbesserung von Produkten. Das ist auch passiert – aber langsamer, als es hätte sein müssen.

Gerade europäische Universitäten sind ja sehr gut in der Grundlagenforschung, aber nicht in der Anwendung. Ist Google auch in diese Wissensfalle getappt?

JAKOB: Ich glaube, das ist ein anderes Phänomen. Es hat mehr mit der Wichtigkeit Googles für unsere Gesellschaft zu tun. Sicherlich sind da Fehler passiert. Aber insgesamt legt Google ein Maß an Sorgfalt an, das zwar verlangsamt wirkt, das aber sein muss.

Die Entwicklungsgeschwindigkeit ist ja wirklich enorm. Eine Gruppe namhafter KI-Fachleute und Unternehmer um Elon Musk hat vor einigen Wochen ein Moratorium für die KI-Entwicklung gefordert. Was halten Sie davon?

JAKOB: Es ist nicht durchsetzbar. Der Zugang zu dieser Technologie verbreitet sich momentan so rasant, dass das gar nicht mehr überprüfbar ist. HANS: Der Aufruf richtet sich zum Teil auch gegen etwas Falsches. Das Ziel ist auch noch falsch.

Was meinen Sie?

HANS: Es könnte Weiterentwicklungen geben, bei denen man wirklich sehr vorsichtig sein muss. Aber dieser Aufruf richtet sich auch gegen Formen der KI, die qua ihrer Funktionsweise überhaupt nicht böse und schlecht werden können. Die Modelle haben im Moment keine Autonomie. Wenn ich nichts eingabe, passiert auch nichts. Die Modelle denken, nicht weiter und füttern sich auch nicht selbst. Der Aufruf ist daher viel zu weit gefasst. Das wäre so, als würde die Biologie sagen: Jetzt hören wir überhaupt mal auf, an Genetik zu arbeiten, nur weil wir nicht Klonen wollen.

Wo könnten denn die tatsächlichen Gefahren in der Zukunft liegen?

HANS: Forscher könnten den Modellen Autonomie verleihen. Wenn man dann die Maschine nicht isoliert hält,

sondern sie mit anderen Dingen interagieren lässt, sie Entscheidungen treffen lässt, dann ist Vorsicht geboten. Aber so, wie der Aufruf formuliert war – das ist einfach nur Alarmschuss. JAKOB: Nicht nur das: Es ist sogar noch etwas schlimmer. Es wird sehr viel Zeit damit zugebracht, zu überlegen, was passieren könnte, wenn solche Modelle einen eigenen inneren Antrieb entwickeln. Was könnten die dann alles kaputt machen? Das ist auch etwas, worüber man nachdenken muss. Aber kurz- und mittelfristig ist es viel problematischer, was Menschen mit diesen Werkzeugen anstellen, deren Intentionen nicht die besten sind. Nicht nur Kriminelle, sondern auch Nationalstaaten, die andere Werte als wir haben.

Was ist dagegen zu tun?

JAKOB: Tatsächlich ist die Weiterentwicklung der Technologie die einzige Hoffnung, um solche Risiken zu kontrollieren.

Wann haben Sie entschieden, in die Spuren Ihres Vaters zu treten?

JAKOB: 2006. Die haben mir dann unterschiedliche Projekte zur Auswahl gestellt, an denen ich arbeiten könnte. Da hatte für mich die spannendsten algorithmischen Probleme maschinelle Übersetzung. Ich wusste aber auch, wenn ich das mache, rücke ich thematisch wieder sehr eng an das heran, was mein Vater machte. War doch einer seiner ehemaligen Doktoranden auch in der Gruppe. Und ich wusste: Wenn ich da reingehe, ist klar, dass ich da erst mal so etwas wie der Sohn von Hans sein werde.

Wann war das?

JAKOB: 2006. Die haben mir dann unterschiedliche Projekte zur Auswahl gestellt, an denen ich arbeiten könnte. Da hatte für mich die spannendsten algorithmischen Probleme maschinelle Übersetzung. Ich wusste aber auch, wenn ich das mache, rücke ich thematisch wieder sehr eng an das heran, was mein Vater machte. War doch einer seiner ehemaligen Doktoranden auch in der Gruppe. Und ich wusste: Wenn ich da reingehe, ist klar, dass ich da erst mal so etwas wie der Sohn von Hans sein werde.

Und da haben Sie sich ausgetauscht?

HANS: Klar. Jakob war bei seiner ersten Gründung noch nicht achtzehn, und da musste ich für ihn die Geschlechterrolle

Sie kamen also inhaltlich und thematisch wieder zusammen.

HANS: Genau. Mathematisch fundierte Algorithmen wurden in der Sprachverarbeitung geradezu entscheidend – und wir kamen dadurch wieder zusammen.

Damit hat sich Ihre Absetzbewegung aus der Jugend ...

JAKOB: ... ganz offenbar verändert. Das Interessante dabei ist, und das haben wir so bislang noch nie gesagt, dass die Transformer, also das T in GPT, von der Linguistik inspiriert worden ist.

Quasi ein Ergebnis Ihrer Gespräche?

JAKOB: Vielleicht hätte das Ganze auch nicht funktioniert, wenn wir nicht schon sehr früh unsere Unterhaltungen am Abendbrotisch dazu gehabt hätten.

HANS: Dabei ging es vor allem darum, Dinge zusammenzubringen, die nicht nebeneinanderstehen, aber in diesen speziellen Anwendungen eigentlich zusammengehören. Wir mussten also die rein mathematischen Verfahren und die linguistischen Verfahren zusammenbringen. Und das kann das T in GPT, also der Transformer. Der kann genau das: Sachen über den Attention-Mechanismus zusammenbringen, die nicht nebeneinanderstehen, aber thematisch miteinander verwandt sind.

Das konnten bisherige Sprachmodelle nicht?

HANS: Richtig, die konnten das nicht. Erst durch die Transformer-Architektur sind die jetzigen Sprachmodelle so mächtig geworden, dass sie das eine mit dem anderen verbinden.

Wir haben mit Ihnen ja auch zwei Gründer, von denen der eine sein Start-up in den USA und der andere sein Start-up in Deutschland ins Leben gerufen hat.

HANS: Genau genommen hat Jakob seine Firma in den USA und in Berlin gegründet; und ich habe in Peking und Berlin gegründet. Berlin ist faktisch unser Schnittpunkt.

Und da haben Sie sich ausgetauscht?

HANS: Klar. Jakob war bei seiner ersten Gründung noch nicht achtzehn, und da musste ich für ihn die Geschlechterrolle

spielen. Aber eigentlich war er sein eigener Gründer.

Vorwurf für hinauswollen, ist die Standortorthematik.

HANS: Da sind wir hier in Europa sehr hinterher. Da machen wir uns beide große Sorgen.

Das haben Sie ja auch gegenüber der deutschen Regierung angesprochen.

HANS: Und auch gegenüber der deutschen Industrie. Die Regierung hat gesagt: Da müssen wir natürlich sofort was machen. Für dieses Jahr aber hat das Parlament den Haushalt schon gemacht. Aber vielleicht gleich im nächsten Jahr – da müssen wir mal gucken. Wenn die Regierung jedoch sagt, wir müssen etwas sofort machen, ist das für uns immer noch sehr weit weg.

Was meinen Sie?

HANS: Dass wir nicht unbedingt das nachmachen müssen, was die Amerikaner vorgemacht haben.

Wir sollten also Modelle primär für die Wirtschaft entwickeln?

HANS: Ja, auf der Konsumentenseite haben wir als Europa oder gar Deutschland keine Chance mehr. Im Industriegeschäft sieht das ganz anders aus.

Das heißt?

HANS: Ich komme gerade aus China zurück. Das Tempo dort ist atemberaubend. Dort geht alles Zack, zack, zack. Die sind einfach so viel schneller.

Sie glauben aber daran, den Vorsprung der Amerikaner aufhalten zu können, sonst würden Sie hier nicht das machen, was Sie machen, oder?

HANS: Wenn wir jetzt behaupten würden, wir werden die Nummer eins – dann wäre das ein Hirngespinnst. Da haben die Amerikaner viel zu viel Masse, viel zu viele Forschungseinrichtungen.

Aber?

HANS: Aber: Es gibt viele Gebiete, in die viele dieser Modelle jetzt erst reinkommen. Die Baidus, Googles, Open AIs und Microsofts haben ja eigentlich Modelle, die hauptsächlich für den Gebrauch durch Endkonsumenten sind. Europa hat immer eine Chance, wenn es um Anwendungen geht, die sich um Unternehmen und Industrien drehen. Und lassen Sie mich auch sagen, diese Modelle von Google, Microsoft und Open AI haben ja auch Defizite.

Welche?

HANS: Nehmen Sie nur die europäischen Sprachen. Davon gibt es mehr als zwei-

hundert. Ungefähr die Hälfte ist in den großen Sprachmodellen der Techniesen nicht gut oder gar nicht abgedeckt.

Warum?

HANS: Viele europäische Sprachen werden von relativ wenigen Menschen gesprochen. Sie werden daher einfach als zu klein eingestuft. Die Modelle sind eben nach anderen Kriterien gemacht. Das heißt ja nicht, dass wir die gleichen Modelle machen und den gleichen Kriterien folgen müssen.

Was meinen Sie?

HANS: Dass wir nicht unbedingt das nachmachen müssen, was die Amerikaner vorgemacht haben.

Wir sollten also Modelle primär für die Wirtschaft entwickeln?

HANS: Ja, auf der Konsumentenseite haben wir als Europa oder gar Deutschland keine Chance mehr. Im Industriegeschäft sieht das ganz anders aus.

Das heißt?

HANS: In der wirtschaftlichen Umsetzung ist das große Geschäft dieser Modelle vielleicht gar nicht im Konsumentenbereich, sondern im B2B-Geschäft, im verarbeitenden Gewerbe zu sehen. Hier öffnen sich völlig neue Anwendungen, ja Welten.

Und was machen hiesige Unternehmen?

HANS: Auf der Industrie- und der

Also braucht Europa so etwas wie ein Google für die Industrie?

JAKOB: Europa braucht vielleicht noch etwas ganz anderes.

Was?

JAKOB: Etwas, das für die Industrie extreme Vorteile hätte: den Zugang zu solchen Modellen als allgemeines Gut zu betrachten. Diese Modelle sind trainiert auf Daten, die erzeugt werden durch die Allgemeinheit. Wenn wir Produktivität maximieren wollen, müssen wir die Modelle also so etwas wie Wasser oder Elektrizität sehen, also als Utility, als eine Art Allgemeingut.

HANS: Wir haben mit Geldern der EU in einem europäischen Network of Excellence von 60 Forschungszentren in 34 Ländern über Jahre hinweg ganz viele hochqualitative Daten in verschiedenen europäischen Sprachen gesammelt. Daten von einer Qualität, wie sie wahrscheinlich nicht die Amerikaner haben. Das heißt: Wir könnten schon ganz vorn mitspielen. Aber wir müssen uns ranhalten.

JAKOB: Genau. Und um es so richtig kompliziert zu machen, werfe ich das Thema des Urheberrechts in die Runde. Denn so, wie wir es kennen, ist es vorbei. Dabei ist es noch völlig unklar, was an dessen Stelle treten wird.

Was könnte es sein?

JAKOB: Daten müssen ja erstmals erzeugt werden, um eine KI und die Modelle damit zu trainieren. Das muss berücksichtigt werden. Ein Weg wäre es, zu sagen: Solche Daten sind in der öffentlichen Hand.

China hat ja vor drei Jahren schon gesagt: Daten seien als Produktionsfaktor zu betrachten – so wie Kapital, Arbeit, Grund und Boden.

HANS: Genau. Da passiert eine spannende Geschichte. Gerade erst hat die Regierung dort eine strenge Datenverordnung erlassen. Das macht es für KI-Firmen noch mal schwieriger, Daten zu erhalten, um ihre Modelle zu bauen. Im Gegensatz setzt sich der Staat mit diesen Ansprüchen in die Pflicht, die nötigen Daten zu liefern.

Aufseiten der Anbieter solcher Modelle sind wir ja auch nicht ganz unbeliebt. Denken wir nur an Unternehmen wie die Heidelberger Leaph Alpha.

HANS: Genau. Die gefallen mir gut. Das ist auch sehr mutig, was die da machen. Die gehen in die richtige Richtung.

JAKOB: Die Idee, das Konzept digitale Konsumprodukte mit dem Vertriebskanal Internet auf die Industrie zu übertragen, funktioniert scheinbar nicht so gut, wie man sich das seitens der Techniesen einst ausgerechnet hatte. Jedenfalls bakt es bei denen auf diesem Gebiet ganz schön.

Lassen sich die Daten nicht auf Datenmärkten beschaffen?

HANS: Das schon. Aber je nachdem, mit welchen Daten ein System trainiert wird, fallen die Antworten aus. Ich habe die gleiche Frage einem amerikanischen und einem chinesischen System gestellt – und zwei verschiedene Antworten erhalten.

Was war die Frage, und was waren die Antworten?

HANS: Die Frage war: Was kann man machen, wenn der Aufnahmeartrag für die kommunistische Partei Chinas abgelehnt wird?

Und die Antwort?

HANS: Das chinesische System zählte genau auf, was man alles machen kann. Das amerikanische System sagte nur, das ist halt ein totalitäres System, da könne man nichts machen.

JAKOB: Es ist ja nicht verwunderlich, dass es da verschiedene Antworten gibt.

Wäre das ein Modell für Europa, dass Daten ein öffentliches Gut sind?

HANS: Ja, auf alle Fälle. JAKOB: Das denke ich auch. Das wäre für ganz Europa auch nicht so teuer. So könnten auch mittelständische Unternehmen den potentiellen Nutzen von KI viel besser einschätzen.

Wer könnte die Entwicklung in diese Richtung anstoßen?

HANS: Wir haben uns ja stets dafür ausgesprochen, dass eines der großen Modelle Open Source ist. Das muss einfach als Commodity da sein, dann kann Europa auch wieder etwas machen.

JAKOB: Das wird aber nicht helfen, wenn die Entwicklungstand so weit hinterher ist, dass man das Geld dann doch wieder den Amerikanern oder Chinesen gibt.

Sehen Sie durch die Spannungen in der globalen politischen Lage, ob es damit einhergehenden Handelsaktionen bis hin zum Chip-Krieg einen neuen KI-Winter heraufziehen?

JAKOB: Wir sehen ja heute schon eine Verlangsamung der Entwicklung. Die Exportbeschränkungen etwa für Maschinen der Entwicklung generativer KI.

Die KI-Familie

Hans Uszkoreit ist Informatiker und hat sich auf die Themenbereiche Computerlinguistik und Künstliche Intelligenz (KI) spezialisiert. Seit 1997 leitet er das Deutsche Forschungszentrum für Künstliche Intelligenz. 1990 in Rostock geboren, wuchs Uszkoreit in der DDR auf. Als oppositioneller Jugendlicher musste er ins Gefängnis. Später floh er nach Westberlin. Dort war er an der Entstehung des bekannten Stadtmagazins „Zitty“ beteiligt, bevor er später Linguistik und Informatik studierte. Uszkoreit ist maßgeblich an der Initiative LEAM des KI-Bundesverbandes beteiligt, die gegenüber Politik und Wirtschaft für die Entwicklung europäischer großer Sprachmodelle wirbt. Der Wissenschaftler gründete mehrere Start-ups, zuletzt das KI-Unternehmen Nyonio. Nyonio will generative KI für Europa entwickeln – Uszkoreit war es leid, auf die Politik zu warten. Im Gründungsteam von Nyonio ist auch seine Frau Feiyu Xu, die dafür ihren Job als KI-Verantwortliche des Softwarekonzerns SAP aufgab.

Jakob Uszkoreit ist der Sohn von Hans Uszkoreit. Er studierte Informatik und Mathematik in Berlin. Nach dem Studium begann Jakob Uszkoreit seine Karriere bei Google und forschte dort später an neuronalen Netzen. Er war Mitautor des viel beachteten Google-Papiers „Attention is all you need“, welches die Grundlagen für die Transformertechnologie legte. Auf dieser Technologie basiert unter anderem auch ChatGPT. 2021 verließ er Google, um Inceptive sein eigenes Start-up zu gründen. Inceptive forscht an der Schnittstelle zwischen Biologie und KI. Jakobs Schwester Lena arbeitet immer noch für Google und ist dort in leitender Position unter anderem für das Testen von Wearables verantwortlich.

KI-Kenner:
Hans Uszkoreit (rechts) und sein Sohn Jakob halten die Panik vor Künstlicher Intelligenz für übertrieben.

Foto: Hannes Jang

gen daraus lassen sich heute noch gar nicht absehen.

Wäre das nicht gut für Europa?

HANS: Wieso?

Es gibt uns Zeit. Denn wenn Europa schon aufholen will, wäre es doch gut, wenn vorn nicht mehr ganz so viel Tempo gemacht wird.

HANS: Richtig ist: Europa muss schneller werden. Es hat aber das Problem, dass es viel Zeit in langwierigen Prozessen der Verwaltung, bei den Abstimmungen und dergleichen verliert.

Und wie kommen wir da wieder raus?

JAKOB: Ich glaube an die Philantropie.

An Philanthropie?

JAKOB: Die Politik ist zu langsam, und die Wirtschaft ist in ihrer Struktur zu heterogen, um zur gleichen Zeit an einem Strang in die gleiche Richtung zu ziehen. Die beiden kriegen es allein nicht hin.

Also?

JAKOB: Es gibt aber genug sehr wohlhabende Europäer, die technisch kompetent sind, Interesse haben und sich auch engagieren wollen.

HANS: Wir haben gesagt, gut durchgerechnet brauchen wir für einen ersten Schritt 350 bis 500 Millionen Euro. Das ist nicht zu viel.

Aber auch nicht gerade wenig?

HANS: Okay. Wie viel soll der neue Aufbau an Kanzleramt kosten? 750 Millionen Euro! Wie viel kosten 50 Kilometer von den Autobahnen, die jetzt neu gebaut werden? Auch so um die 350 Millionen Euro! Man darf nicht darüber nachdenken – da wird man ja ganz trübsinnig. Wenn man sieht, wie die in der Regierung alle nicken und sagen: Ja, ist uns ganz wichtig, toll, dass Sie uns das so gut erklärt haben; das ist die größte Entwicklung seit der Elektrizität. Und dann kommt nichts. Also, da fehlen mir die Worte. Deshalb starten wir jetzt mit unserem frisch gegründeten KI-Unternehmen Nyonio erst mal privat mit der Entwicklung generativer KI.

Das Gespräch führte: **Maximilian Sachsse und Stephan Finsterbusch**