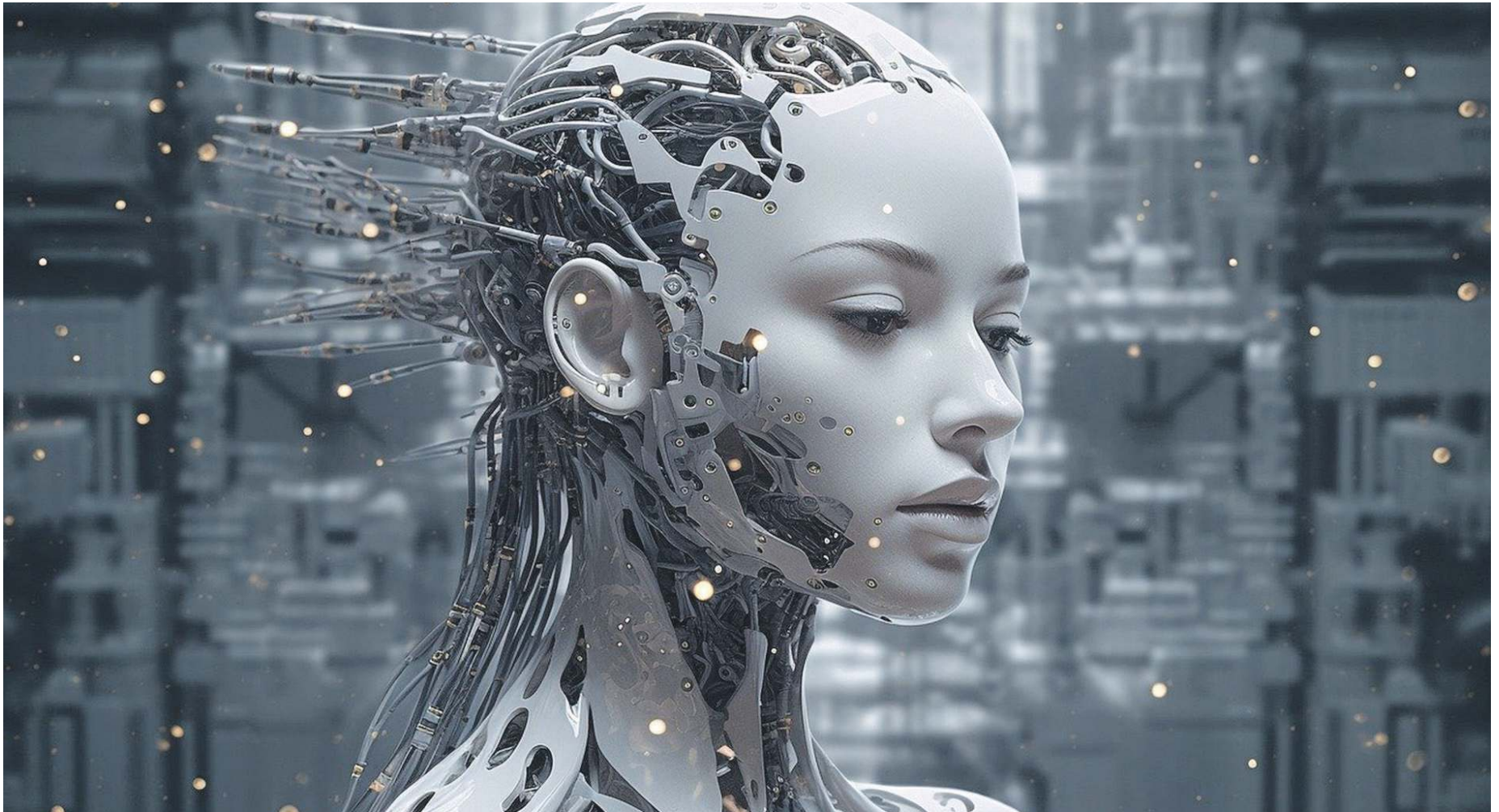




Das Potenzial der Künstlichen Intelligenz wird von bekannten Experten über das der Atombombe gestellt: Segen und Fluch liegen eng beieinander.

Foto: Pixabay

KI macht Betrug echter

Mittels Künstlicher Intelligenz haben es Betrüger jetzt schon kinderleicht – Experten warnen und fordern Regulierungen

ds. Die Künstliche Intelligenz (kurz: KI) ist da und damit ihr enormes, noch unvorhersehbares Potenzial, das auf unsere Gesellschaft einwirken wird. Ob Segen oder Fluch liegt letztendlich in Menschenhand. Zur Schattenseite der KI gehören unter anderem ausgeklügelte und gefährlichere Betrugsmaschinen, kinderleicht von Menschen mit kriminellen Absichten zu erstellen, auf die Privatpersonen sowie Unternehmen gleichermaßen immer häufiger hereinfallen.

Das Problem dabei: Echt oder falsch sind fast nicht mehr zu unterscheiden. So wurde kürzlich von einem chinesischen Geschäftsmann berichtet, der durch einen KI-Betrug 4,3 Millionen Yuan, etwa 567.000 Euro, verlor. Der Betrüger gab sich mit Hilfe von Künstlicher Intelligenz in einem Videoanruf sowohl optisch als auch stimmlich als ein enger Freund des Opfers aus.

Echt oder falsch?

Echt oder falsch mit Hilfe der neuen Technologien ist kaum noch zu

unterscheiden – in Zukunft wahrscheinlich gar nicht mehr. Aufhorchen lässt die Aussage von Sam Altman, CEO von Open AI, dem KI-Unternehmen hinter dem medial bekannten KI-basierten Chatbot ChatGPT. Er ist überzeugt: der Mensch müsse künftig mittels digitaler Identität nachweisen, dass er keine KI ist.

Hielt man die KI-Technologie noch bis kürzlich für etwas „Ungreifbares in der Ferne“, hat sie nun spätestens mit ChatGPT, einer Sprach-KI, die mit ihren Nutzern in natürlicher menschlicher Sprache interagieren kann, die Welt „im Sturm erobert“. Und ChatGPT ist nur ein Teil vom Ganzen, das gerade erst begonnen hat. Der KI-Geist ist aus der Flasche und wird mancherorts gefeiert und bejubelt wie ein neuer Weltstar.

Andere wiederum lassen Skepsis walten. Gerade in den letzten Wochen gab es immer wieder Experten und Schlagzeilen, die zur Vorsicht mahnten und eine stärkere Regulierung seitens der Staaten forderten, darunter gar der milliar-

denschwere Tech-Mogul Elon Musk – übrigens einer der Mitgründer von Open AI –, der davor warnte, dass unkontrollierte Künstliche Intelligenz eine Gefahr für die Gesellschaft darstelle.

Erst im Mai verließ der angesehene „Urvater der KI“, Geoffrey Hinton, Google, wo er in den letzten zehn Jahren im Bereich „Deep Learning“ (einem Teilgebiet des maschinellen Lernens, welches

sich auf künstliche neuronale Netze und große Datenmengen fokussiert) arbeitete, um ungehindert aufzurütteln. Nach seiner Meinung, die er gegenüber der New York Times (NYT) äußerte, wird es nicht möglich sein, „böse Akteure daran zu hindern, die Technologie für böse Dinge zu nutzen“.

Die Gefahr: Das Internet kann mit gefakten, also gefälschten Bildern, Videos und Texten geflutet

werden und der durchschnittliche User würde nicht mehr in der Lage sein zu erkennen, was noch wahr und was falsch ist. Im Mai titelte NTV: „Fake-Foto von Pentagon-Explosion sorgt für Wirbel. [...] Das vermutlich von einer Künstlichen Intelligenz (KI) erzeugte Bild wurde von mehreren Nutzerkonten in den Internetdiensten geteilt, sodass sich das US-Verteidigungsministerium zu einer Stellungnahme gezwungen sah.“

KI-Chatbots sind „ziemlich beängstigend“, so Hinton. Soweit er das beurteilen könne, seien die Chatbots im Moment noch nicht intelligenter als die Menschen, aber er denke, sie könnten es bald sein, indem sie den Informationsgehalt des menschlichen Gehirns übertreffen. Die künstliche Intelligenz arbeitet mit künstlichen neuronalen Netzen, die dem menschlichen Gehirn nachempfunden sind und entsprechend so lernen und Informationen verarbeiten.

Menschen sind „biologische Systeme“, KI dagegen ist digital. Der große Unterschied besteht laut



In der Mitte: KI-Pionier Geoffrey Hinton. Foto: Steve Jurvetson, Wikimedia

Hinton darin, dass es bei digitalen Systemen viele Kopien derselben Gewichtung, desselben Modells in der Welt geben kann. „Und all diese Kopien können getrennt voneinander lernen, aber ihr Wissen sofort weitergeben. Das ist so, als hätte man 10.000 Leute, und immer, wenn eine Person etwas gelernt hat, wissen es automatisch alle. Und so können diese Chatbots so viel mehr wissen als eine einzelne Person.“

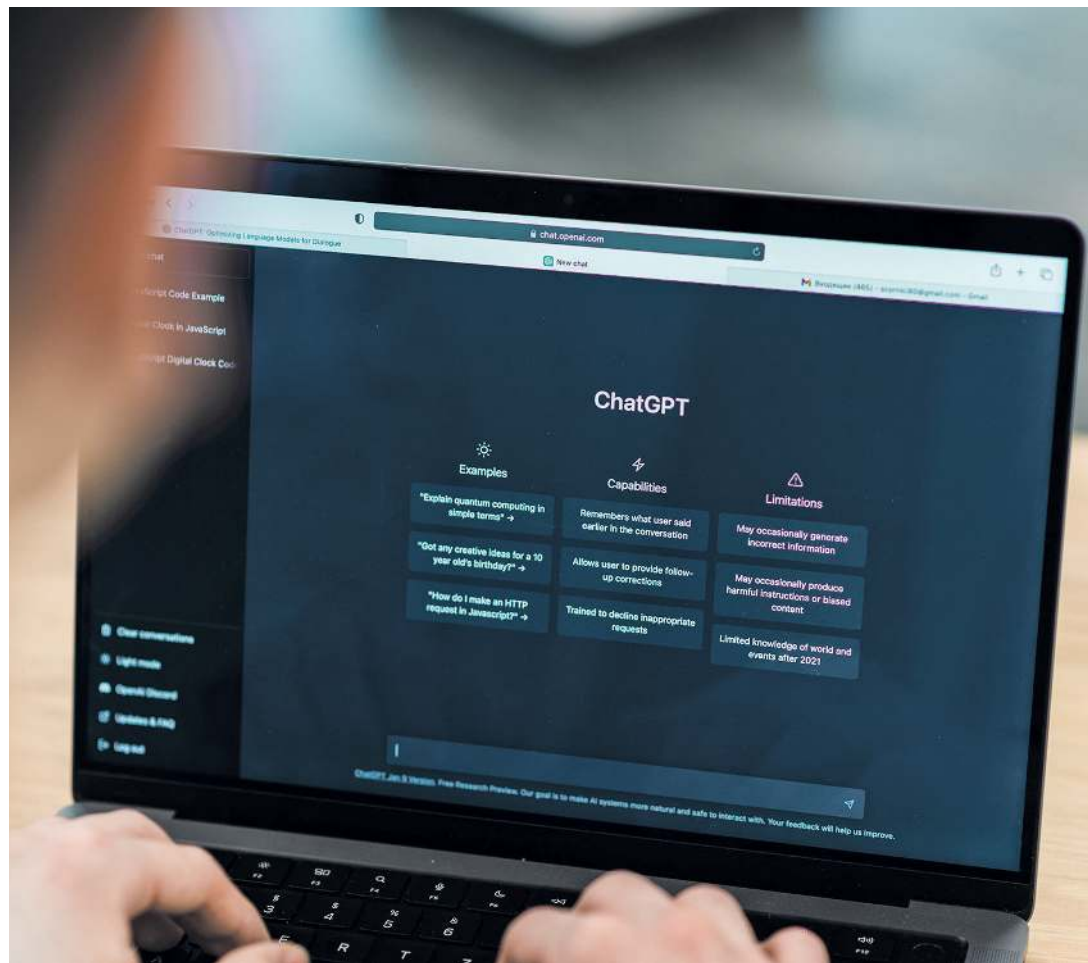
Futuristen beschreiben – in einem Video der Hackergruppe Anonymous Schweiz anschaulich erklärt – den Moment, in dem die KI die Intelligenz und die Schnelligkeit des menschlichen Gehirns übertrifft und sich dabei weiter kontinuierlich selbst verbessert, als Technologische Singularität. Ab diesem Zeitpunkt wird es automatisch zu einem Sprung des technologischen Fortschritts kommen. Dann wird dieser von Menschen geschaffene Computer in der Lage sein, jegliches Problem – ob technisch, wissenschaftlich oder philosophisch –, an dem Menschen schon lange arbeiten, besser und schneller zu lösen, was Chancen und Risiken birgt. Und der technologische Fortschritt schreitet rasant schnell voran. So soll der Fortschritt vom Jahr 2000 bis 2020 genauso groß gewesen sein wie der des gesamten 20. Jahrhunderts.

Kein Wunder also, dass bei einer so mächtigen Erfindung ein regelrechter Wettlauf zwischen den Tech-Giganten stattfindet. So sagte Elon Musk gegenüber Fox News, dass Larry Page von Google mit der Artificial General Intelligence (Allgemeine künstliche Intelligenz, kurz AGI) eine „digitale Superintelligenz, im Grunde einen ‚digitalen Gott‘ schaffen will, von der viele Silicon-Valley-Ingenieure glauben, dass er der Schlüssel für das Eintreten der Technologischen Singularität oder der Intelligenzexplosion ist.“

Neben der Aussage „Wenn die KI intelligenter wird als der Mensch, wird die Zukunft unvorhersehbar“ gibt es bereits konkretere und reale Beispiele von Schattenseiten wie Phishing, Identitätsdiebstahl, groß angelegte Desinformationen (Fake News) und Cyberangriffe. Natürlich versuchen Tech-Unternehmen Missbräuche ihrer KI-Erfindungen mit Hinweisen und Einschränkungen zu verhindern, aber dennoch scheint es immer wieder Möglichkeiten zu geben, die zu umgehen.

Nicht nur geklonte Stimmen

So entwickelt sich die Zunahme von Betrugsmaschinen mittels KI mittlerweile zu einem globalen Trend. Werden Fotos, Videos oder Audio-Dateien mit Hilfe von Künstlicher Intelligenz absichtlich verändert, spricht man von Deepfake.



Der am rasantesten wachsende Dienst aller Zeiten: die Sprach-KI ChatGPT. Foto: frimufilms, Freepik

Beachtung findet bei Betrügern offenbar der alte „Enkeltrick“, der nun eine neue Stufe erreicht hat: Dank KI können täuschend echte Stimmen generiert, sogar geklont und eingesetzt werden. Ausschlaggebend dabei ist, dass die Sprach-KI zum einen die passenden Emotionen im Text erkennen muss und zum anderen muss sie in der Lage sein, diese Emotionen durch die Stimme richtig auszudrücken. Klingt kompliziert, aber diese Hürden wurden inzwischen von der KI genommen.

So berichteten spanische Medien über die 73-jährige Ruth Card, die einen Anruf von ihrem vermeintlichen Enkel Brandon erhielt, bei dem es sich aber in Wirklichkeit um Betrüger handelte: „Oma, ich bin im Gefängnis, keine Geldbörse, kein Telefon. Ich brauche

Geld für die Kautions.“ Oder der Fall von Jennifer DeStefanos 15-jähriger Tochter Briana, über den WKYT.com berichtete.

Es war ihre Stimme

Die Tochter befand sich im Ski-Urlaub, als die Mutter plötzlich einen verstörenden Anruf ihrer angeblichen Tochter erhielt. „Mama, ich habe Mist gebaut“. Dann hörte sie eine Männerstimme: „Hören Sie zu. Ich habe Ihre Tochter. Ich sage Ihnen, wie es ablaufen wird...“ Während der Lösegeldforderung nahm sie im Hintergrund die schluchzende Stimme ihrer Tochter wahr: „Hilf mir, Mama. Bitte hilf mir.“ Glücklicherweise wurde das Vorhaben schnell vereitelt. Dennoch, während des Anrufs bezweifelte DeStefano keinen Moment, dass es sich am Telefon um

ihre Tochter handelte. „Es war genau ihre Stimme. Es war ihr Tonfall. Es war die Art, wie sie geweint hätte“, sagte sie. „Das ist der verrückte Teil, der mich wirklich bis ins Mark getroffen hat.“

Wie ist das möglich? Nur wenige Sekunden einer „Sprachprobe“ – ein kurzes Gespräch mit einem unbekannten Anrufer, Videos bei YouTube, TikTok und Co., Sprachnachrichten und so weiter – genügen, um diese später für kriminelle Zwecke missbrauchen zu können. Von daher ist Vorsicht angeraten. Nie vorschnell handeln!

Kommt einem ein Anruf seltsam vor, kann man beispielsweise nach sehr persönlichen Dingen fragen, die ein Betrüger nicht kennen kann. Oder aber es wird in der Familie ein Passwort ausgemacht, das nur die Familienmitglieder kennen können.

Beim Betrug mit „KI-geklonter Stimme“ bleibt es leider nicht. In anderen Fällen wurde ein Video-Klon einer Person verwendet, um Angehörige reinzulegen. Wie die spanische Zeitung „El Mundo“ berichtet, erhielt ein Mann in China über WeChat einen Videoanruf von seiner angeblichen Frau. Sie bat ihn um 3.600 Euro, weil sie einen Autounfall gehabt habe und die Situation mit dem anderen Fahrer regeln müsse.

Natürlich war es nicht seine Frau. Und obwohl sie ihn von einem anderen Konto aus anrief, tappte der Mann in die Falle, denn das Gesicht, das in dem Videoanruf erschien, gestikulierte und sprach in demselben Tonfall wie seine Frau.

Allerdings handelte es sich um eine KI-Imitation. Offenbar hatten die Betrüger das Paar beobachtet und kannten seine Gewohnheiten. Außerdem betrieb die Frau einen beliebten Kochkanal in einem sozialen Netzwerk, von dem die Betrüger die Aufnahmen ihres Gesichts und ihrer Stimme für die Erstellung des Deepfakes bezogen.

Dass KI in der Lage ist, „Unwahres“ täuschend echt erscheinen zu lassen, birgt Risiken in zweierlei Hinsicht: Durch Deepfake können Menschen leicht betrogen werden. Das heißt, etwas, das nicht real ist, erscheint real. Deshalb werden Menschen beginnen, Bild- und Tonaufnahmen immer weniger Glauben zu schenken. Das heißt: Auch wenn etwas real ist, wird man glauben, es könnte unreal sein. Was wird allein dieser Aspekt aus einer Gesellschaft machen?

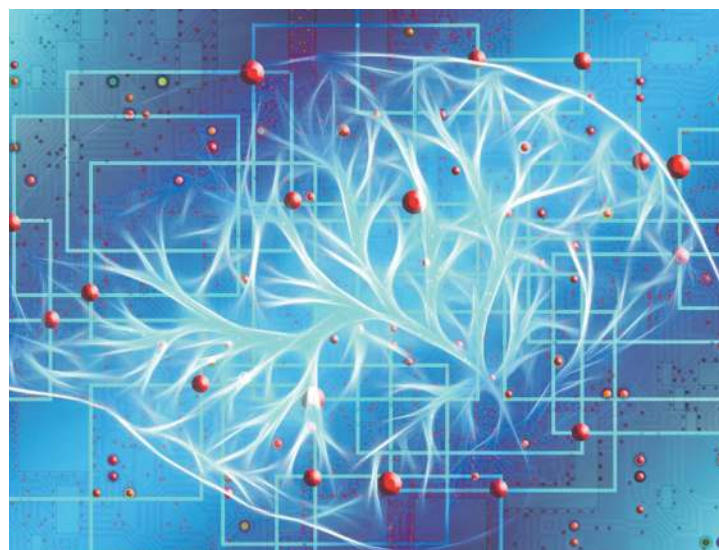
KIs wie ChatGPT können auch programmieren. Das kann im Guten zur Vorbeugung als auch im Schlechten als Angriff genutzt werden. „Für einen potenziellen Kriminellen mit geringen technischen Kenntnissen ist dies eine unschätzbare Ressource, um einen bösartigen Code zu erstellen“, so Europol, die Strafverfolgungsbehörde der Europäischen Union.

Regulierungen sind schwierig

Insofern wundert es nicht, dass Regulierungen beginnen, über die Regulierung der KI zu sprechen. So haben sich die G7-Staaten bei ihrem Treffen im Mai in Hiroshima darauf geeinigt, bis Ende des Jahres neue Standards vorzuschlagen. Auch die Europäische Union arbeitet an einer Regulierung, dem „Künstlichen-Intelligenz-Gesetz“, kurz „AI-Act“. Wahrscheinlich werden KI-Anwendungen in vier Risiko-Kategorien eingeteilt: minimales, begrenztes, hohes und inakzeptables Risiko. Das Unterfangen ist kompliziert, wissen nicht einmal die Experten selbst in manchen Situationen ihre KI zu deuten – wie sollen es dann andere tun?

So verriet Google-CEO Sundar Pichai in einem Interview mit CBS News: „Es gibt einen Aspekt, den wir alle, die wir in diesem Bereich arbeiten, als Blackbox bezeichnen. Man versteht nicht ganz, und man weiß nicht genau, warum sie (die KI) dies gesagt hat oder warum sie einen Fehler gemacht hat.“

Die EU-Regeln zur KI sollen Ende des Jahres stehen und dann 2025 in Kraft treten. Whistleblower und IT-Experte Edward Snowden twitterte diesbezüglich: „Die Institutionen versuchen, die künstliche Intelligenz zu beherrschen und zu lenken, die ihrerseits das Verhalten der menschlichen Intelligenz formen und kontrollieren wird – ein System der doppelten Fesselung.“



Dem menschlichen Gehirn nachempfunden: Künstliche neuronale Netze für die KI. Foto: Gerd Altmann, Pixabay