



21.10.2025

Transkript

„Einfluss künstlicher Intelligenz auf den Forschungsprozess“

Expertinnen und Experten auf dem Podium

- ▶ **Prof. Dr. Clara Schoeder**
Juniorprofessorin für die Entwicklung von immuntherapeutischen Wirkstoffen, Institut für Wirkstoffentwicklung, Universität Leipzig
- ▶ **Prof. Dr. Mario Krenn**
Professor für Maschinelles Lernen in der Wissenschaft, Eberhard Karls Universität Tübingen
- ▶ **Prof. Dr. Florian Boge**
Professor für Wissenschaftsphilosophie mit dem Schwerpunkt Künstliche Intelligenz, Technische Universität Dortmund
- ▶ **Samantha Hofmann**
Redakteurin für Digitales und Technologie, Science Media Center Germany und Moderatorin dieser Veranstaltung

Mitschnitt

- ▶ Einen Audio- und Videomitschnitt finden Sie unter:
<https://sciencemediacenter.de/angebote/einfluss-kuenstlicher-intelligenz-auf-den-forschungsprozess-25187>



Transkript

Moderatorin [00:00:00]

Ja, herzlich willkommen, liebe Journalistinnen und Journalisten, hier zum Press Briefing des Science Media Center zum Thema Einfluss künstlicher Intelligenz auf den Forschungsprozess. Mein Name ist Samantha Hofmann, ich bin Redakteurin hier beim Science Media Center für Digitales und Technologie und mit mir hier sind drei Forschende, die sich mit dem Thema KI in der Forschung beschäftigen. Vielen Dank, dass Sie sich die Zeit genommen haben, hier heute unsere Fragen zu beantworten. Bevor wir mit dem eigentlichen Briefing beginnen, noch ein paar organisatorische Dinge. Liebe Journalistinnen und Journalisten, für Ihre Fragen nutzen Sie bitte das F von Zoom. Ein Kollege von mir sammelt dann die Fragen, postet sie mir in ein anderes Dokument und ich stelle sie hier im Gespräch an die Expertinnen und Experten. Nutzen Sie bitte ausschließlich die F für Ihre Fragen und nicht den Zoom-Chat. Das macht das Ganze für uns etwas übersichtlicher. In dem F gibt es außerdem die Möglichkeit, Fragen, die Kolleg:innen von Ihnen gestellt haben, mit einem Daumen nach oben upzuvoten. So können wir sehen, welche Fragen Sie für besonders relevant halten. Und wenn es ganz viele Fragen gibt, können wir dann gegebenenfalls priorisieren, welche Fragen wir stellen. Dieses Meeting wird aufgezeichnet. Das bedeutet, Sie müssen nichts mitschneiden. Wir stellen ein Transkript, eine Audio und eine Videodatei nach dem Meeting zur Verfügung. Jetzt zum Thema. KI erhält in immer mehr Lebensbereiche Einzug, unter anderem auch in die Forschung. Deswegen hat das Joint Research Centre der EU eine Analyse dazu veröffentlicht, wie KI den wissenschaftlichen Prozess beeinflusst. Morgen findet außerdem eine Konferenz von der Universität Stanford statt, bei der KI als Autor fungiert und die Paper geschrieben beziehungsweise einen wichtigen Teil des wissenschaftlichen Prozesses mit begleitet hat. Die Links zu der Analyse vom Joint Research Centre und zu der Konferenz finden Sie in der Einladungsmail zu diesem Press Briefing. Wir sehen also, sowohl Politik als auch Forschung finden KI in der Wissenschaft wichtig. Deswegen sind wir auch heute hier, um darüber zu sprechen. Und das machen wir natürlich nicht alleine, sondern mit Expert:innen zu dem Thema. Da haben wir als ersten hier Mario Krenn. Er ist Professor für maschinelles Lernen in der Wissenschaft an der Eberhard Karls Universität Tübingen. Dort leitet er das Artificial Scientist Lab. Er untersucht, wie KI im wissenschaftlichen Prozess sinnvoll eingesetzt werden kann und hat einen Hintergrund in Quantenphysik, beschäftigt sich aber auch mit KI-Anwendungen in anderen Wissenschaftsbereichen. Danke, dass Sie hier sind. Als zweite Expertin haben wir Clara Schoeder. Sie ist Juniorprofessorin für die Entwicklung von immuntherapeutischen Wirkstoffen an der Universität Leipzig. Ihr Fokus liegt dabei auf Proteindesign mittels Computertechniken wie zum Beispiel AlphaFold. Und die Runde abschließt Florian Boge. Er ist Juniorprofessor für Wissenschaftsphilosophie mit dem Schwerpunkt künstliche Intelligenz an der Technischen Universität Dortmund. Da beschäftigt er sich unter anderem mit Auswirkungen von KI auf wissenschaftliches Verstehen. Ich habe zunächst einmal eine Eingangsfrage für jeden von Ihnen vorbereitet und würde gerne direkt mit Ihnen anfangen, Herr Krenn. Die Frage: Die Analyse des JRC zu KI in der Wissenschaft betrachtet Einsatzmöglichkeiten in allen Stufen des Forschungsprozesses. Für welche Schritte, zum Beispiel Literaturrecherche, Hypothesengenerierung, Datenauswertung, spielt KI aktuell die größte Rolle und wie wird sie konkret eingesetzt?

Mario Krenn [00:03:15]

Danke schön. Danke für Ihre Einladung. Danke für die Frage. Künstliche Intelligenz wird jetzt schon seit mehreren Jahren in der Wissenschaft erfolgreich eingesetzt, vor allem in Datenanalysen. Das heißt quasi ein Feld, das prädestiniert ist für die Anwendung von künstlicher Intelligenz. Ein ganz konkretes Beispiel wäre hier zum Beispiel AlphaFold, worüber wir im Zuge der Diskussion sicher noch sprechen. Das ist ein Programm, das eine der größten Fragen in der Biochemie gelöst hat und



wofür auch der Nobelpreis voriges Jahr in Chemie vergeben wurde. In den letzten, sagen wir, zwei, drei Jahren gibt es immer mehr einen Fokus auf KI, die beiträgt, selbstständig wissenschaftliche Entdeckungen zu machen. Das sind vor allem Systeme, die nicht notwendigerweise frühere und menschliche Daten benutzen, sondern quasi die Daten selbstständig erzeugen können. Da brauchen sie Zugriff auf andere Systeme. Und dort können sie quasi selbstständig neue wissenschaftliche Entdeckungen machen. Das kann man sich zum Beispiel in der Physik vorstellen. Dort können Computerprogramme neue Experimente finden, die menschliche Wissenschaftler so nicht gefunden hätten. Genauso in der Materialwissenschaft kann man sich vorstellen, dass Computerprogramme jetzt neuartige Moleküle finden, die irgendwelche Funktionen haben, die super für gewisse Technologien sind, an die Menschen auch nicht gedacht haben. Genauso kann man sich das vorstellen in der Mathematik, dass Computerprogramme Beweise finden würden, die offene Fragen lösen. In der Mathematik ist die Anwendung noch nicht so weit, dass die Programme tatsächlich im Forschungsbereich schon mithelfen. In der Chemie und in der Physik ist das aber tatsächlich schon so. Mein eigenes Forschungsthema ist genau Physik, wo wir versuchen, Computerprogramme zu erzeugen, die uns helfen, neue Physik-Experimente zu designen. Und warum das ein schwieriges Problem ist, kann man sich so vorstellen: Die Möglichkeit an allen Experimenten, die man hypothetisch bauen könnte im Labor, ist absolut gewaltig, riesig. Man kann sich vorstellen, wenn man einen Strahlteiler, Laser und andere Komponenten hat und die auf einen optischen Tisch hinstellt. Die Kombination, die man hinstellen könnte, die wächst gigantisch. Wenn man, glaube ich, fünf Bauteile hat, kann man schon mehr als tausend Experimente bauen. Wenn man zehn Bauteile hat, mehr als achtzigtausend Experimente. Das wächst einfach gewaltig. Und in diesem Raum der Möglichkeiten finden sich Experimente, die extrem nützliche Eigenschaften haben. Wie zum Beispiel Mikroskope, die man in der Medizin und Biologie und so weiter benutzen oder Teleskope, mit denen man astrophysikalische Beobachtungen machen oder Gravitationswellendetektoren mit denen man andere astrophysikalische Effekte messen könnte. Das sind alles Experimente, die irgendwo in diesem großen Raum an möglichen Experimenten sind. Und bis jetzt oder bis, sagen wir, vor fünf Jahren, haben diesen Raum vor allem kreative und erfahrene Menschen durchsucht. Und das ist, was sich jetzt genau ändert. Hier werden jetzt Computerprogramme benutzt, um diesen Raum auf eine ganz andere Art und Weise zu durchsuchen. Auf eine automatische Art und Weise, die ganz anders funktioniert als die menschliche Art. Und hier gibt es schon viele Resultate, in denen die Maschinen Experimente gefunden haben, auf die keine Menschen gekommen wären und diese Experimente dann auch in den Laboren gebaut worden sind. Da gibt es Beispiele in der Quantenphysik, meinem Gebiet: Mehrere der Experimente, die unsere Computerprogramme gefunden haben, wurden dann in Laboren gebaut und haben uns geholfen, neue Effekte in der Quantenphysik zu beobachten, zu analysieren. Andere Gebiete sind in der Gravitationswellenphysik. Das sind quasi die sensitivsten Messapparate, die die Menschheit jemals entdeckt hat. Da gibt es jetzt auch Resultate, in denen Computerprogramme noch sensitivere Experimente bauen, an die auch wieder Menschen nicht gedacht hätten. Das Interessante ist, wenn man jetzt solche Lösungen hat, die besser sind als was Menschen gemacht haben, dann bedeutet das, dass in diesen Lösungen irgendwelche Tricks oder Ideen codiert sind, die wir als Menschen nicht kennen und nicht gekannt haben. Das heißt jetzt: Nicht nur bekommen wir eine Lösung, sondern wir bekommen quasi in dieser Lösung codiert neue Ideen, die man extrahieren und hoffentlich verstehen kann. Also das ist quasi die Richtung, in die es jetzt einen extremen Push gibt, in die vollständig automatische Entdeckung von neuen wissenschaftlichen Erkenntnissen.

Moderatorin [00:08:56]

Mhm. Also KI durchaus ein bisschen mehr als einfach nur ein Werkzeug wie ein Hammer, sondern wirklich was, was auch zum Ideengewinn beitragen kann.



Mario Krenn [00:09:07]

Genau. In, in dem Fall sind die Ideen implizit gefunden aus Computerprogrammen. Das interessiert uns nicht, ob es eine Idee erzeugt, so dass wir die Frage lösen. Aber wir als Menschen können diese Lösungen dann auf eine spezielle Art und Weise analysieren und diese Ideen extrahieren. Der letzte kurze Punkt: Es gibt jetzt noch einen Push, dass man versucht, den gesamten Wissenschaftsprozess zu automatisieren. Man versucht, durch alle wissenschaftlichen Veröffentlichungen, in Millionen von Texten, Fehlstellen oder fehlende Verbindungen zu finden und dort neue Ideen zu erzeugen und diese Wege potenziell automatisch auszuführen und dann potenziell die Resultate automatisch zu generiert. Das wird wahrscheinlich in der Zukunft sehr wichtig werden. Im Moment sind wir noch nicht da, dass vollkommen vom ersten Punkt bis zum letzten Punkt ein Computerprogramm neue Ideen generieren kann.

Moderatorin [00:10:17]

Da kommen wir später auf jeden Fall auch noch mal drauf zu sprechen, was da schon möglich ist und was man sich in Zukunft vorstellen kann. Vielen Dank für die Antwort. Weiter würde ich gerne machen mit Frau Schoeder. Sie haben es eben auch schon mal ganz kurz angesprochen, Herr Krenn – eine der aufsehenerregendsten KI-Anwendungen in der Forschung in den vergangenen Jahren war vermutlich AlphaFold. Frau Schoeder, können Sie für uns einmal zusammenfassen, wo der Fortschritt im Vergleich zu herkömmlichen Methoden im Proteindesign liegt und welche Bedeutung AlphaFold in der Praxis hat?

Clara Schoeder [00:10:47]

Ja, sehr gerne. Vielleicht als Background dazu erstmal: Ich bin hier an der Universität Leipzig in einem Institut, das sich damit beschäftigt, wie wir diese computergestützten Methoden verwenden können, um wirklich neue Wirkstoffe zu generieren. Und uns ist sehr wichtig, dass wir die Algorithmen nicht nur kennen oder beschreiben, sondern auch wirklich anwenden. Also das, was wir tatsächlich machen, ist, wir nehmen diese Algorithmen und versuchen wirklich Wirkstoffkandidaten zu entwickeln, die wir auch charakterisieren. Und das lehrt uns natürlich ganz viel darüber, was diese Algorithmen können. Und zu der ersten Frage: AlphaFold. Warum ist das so wichtig geworden für uns vor allem? Ich glaube, es ist diese sofortige Bereitstellung von Strukturen von Proteinen, die uns die Arbeit so viel erleichtert. Früher konnte man die auch bereits modellieren, aber dieses Modellieren war aufwendig und auch oft fehlerbehaftet. Das heißt nicht, dass AlphaFold das richtig macht, denn es ist genauso gut wie die Datenbasis, auf der es arbeitet. Und das ist eben eine Protein-Datenbank, in der die Proteinstrukturen jahrzehntelang von Wissenschaftlern abgelegt worden sind. Was AlphaFold macht: Es erlaubt es uns, aus einer Aminosäuresequenz – die können wir zum Beispiel aus dem Genom ablesen – Proteinstrukturen vorherzusagen. Diese Strukturen sind sehr wichtig, weil das erlaubt uns manchmal, nicht immer – das ist, würde ich sagen, eine neue Frontier – auf die Funktion zu schließen. Es erlaubt uns aber auch zu überlegen, wie könnte zum Beispiel ein Arzneistoff mit diesem Protein interagieren? Es erlaubt uns, solche Fragen zu stellen wie, wenn ich eine Mutation beobachte, welchen Einfluss hat das auf die Struktur? Auch das ist viel mit computergestützten Vorhersagen verbunden und von denen sind viele nicht unbedingt perfekt, auch heute nicht. Was AlphaFold wirklich revolutioniert hat, ist, dass dieses von Sequenz zur Struktur sehr schnell geht im Vergleich zu früher und dass ich auch nicht mehr große Mengen an Computepower brauche, um auf eine Lösung in vielen zehntausend zu kommen, die vielleicht irgendetwas gebracht hat für den Wissenschaftsprozess. Meistens generiert man fünf Modelle, die sind meistens schon ziemlich gut und dann kann ich direkt loslegen. Also anstatt dass ich zwei, drei Monate damit verbringe, diese Vorarbeiten zu leisten, brauche ich meistens nur noch zwei, drei Stunden, um sofort anfangen zu können. Und das macht natürlich alle Prozesse viel schneller. Für uns wichtig ist das, was AlphaFold macht: diese



Vorhersage zur Struktur. Das kann man umdrehen. Man kann jetzt fragen, wenn ich eine Funktion haben möchte für ein Protein, welche Sequenz bräuchte ich denn dafür? Und das ist das, was wir Proteindesign nennen. Und dafür gab es ja die andere Hälfte des Nobelpreises letzten Jahres, sozusagen für die Umkehrung dieser Methoden. Und da passiert – ganz ähnlich, was Herr Krenn nämlich auch gerade schon angedeutet hat: Wenn der Computer solche Vorschläge für neue Proteine macht, lernen wir dabei ganz viel über unser Verständnis von Proteinen. Wie sich diese Proteine verhalten im Labor – die haben nämlich ganz spannende Eigenschaften – und auch, welche Fehlstellen diese Algorithmen noch haben, weil wir sehen auch, dass es Dinge gibt, die diese Algorithmen nicht abbilden können und wo wir jetzt – das ist unsere neue wissenschaftliche Arbeit – versuchen, das zu optimieren. Und das machen wir zum Beispiel auch hier bei uns im Labor. Wir beobachten also, wenn wir solche De-novo-Proteine designen, welche Eigenschaften die haben.

Moderatorin [00:14:07]

Mhm. Vielen Dank. Auch da hat man gehört, KI spielt auf jeden Fall eine wichtige Rolle im Forschungsprozess und das leitet auch schon ganz gut zur dritten Frage über an Sie, Herr Boge. Und zwar, verändert KI den wissenschaftlichen Erkenntnisprozess – also die Art, wie Wissen generiert wird – fundamental?

Florian Boge [00:14:26]

Ja, vielen Dank auch von meiner Seite für die Einladung und für die Frage. Im Grunde genommen hat Herr Krenn mir schon einiges vorweggenommen. Das ist etwas schade. Ich kann das ja versuchen, wissenschaftsphilosophisch noch mal so ein bisschen aufzuarbeiten. Und zwar gibt es eine gewisse Warte, aus der eigentlich gar nicht so wirklich was Neues hinzukommt. Also so ein KI-Modell kann ja eigentlich immer beschrieben werden als eine relativ komplizierte mathematische Funktion mit sehr vielen freien Größen, die nennt man dann freie Parameter. Und diese freien Größen, die passt man dann in so einem Prozess, den man dann das Lernen des Modells nennt, an. Und ich meine, das ist ja etwas, das hat die Wissenschaft eigentlich schon lange gemacht, schon immer gemacht: Mit Modellen zu arbeiten, bei denen es freie Größen gibt, die man anpassen kann. Jetzt gibt es aber schon ein paar Unterschiede zwischen klassischen Modellen, die man beispielsweise eben im Bereich der Physik für Simulationen, aber auch viel direkter für Berechnungen auf Papier mit dem Stift benutzt hat und KI-Modellen – insbesondere was den Modellierungsprozess angeht. Zum Beispiel geht man bei einem traditionellen Modell etwa hin und guckt sich an: Wie ist denn eigentlich das System, was ich mir anschauen will, aufgebaut? Was hat das für Eigenschaften? Und dann versucht man, die abstrakt mathematisch zu beschreiben, guckt sich die Beziehungen zwischen diesen Eigenschaften an und kann daraus dann Schlussfolgerungen über das System, das einen interessiert, ableiten. Bei der KI macht man das natürlich erst mal gar nicht so, sondern man schaut sich erst mal an: Ist dieses Modell, was ich hier nehme, im Grunde genommen in der Lage, diese Aufgabe, die ich ihm stelle, möglichst zügig, möglichst effizient zu erledigen? Und die Ideen darüber, wie das Zielsystem – sozusagen das, was ich damit untersuchen will – aussieht, die spielen am ehesten dabei eine Rolle, wenn man sich den Lernalgorithmus überlegt. Wenn man zum Beispiel festlegt: Wann macht das System denn einen Fehler? Wann muss ich es denn sozusagen bestrafen, in Führungszeichen, damit es sich anders verhält? Die Zielsysteme kommen aber in einer viel geringeren Weise nur noch vor, als das eben bei klassischen Modellen der Fall ist. Also hier gibt es erst mal schon mal einen grundlegenden Unterschied im Ansatz, wie das Modellieren sozusagen funktioniert. Ein weiterer Unterschied, der aber auch vielleicht irgendwo zumindest fundamentale Konsequenzen hat, ist der, dass ich gleichzeitig auf so ein KI-System mit etwas Fantasie so ein bisschen gucken kann, als wäre das jetzt ein neuer kognitiver Akteur. Ja, als wäre das sozusagen ein weiterer Mitdenkender. Das muss ich vielleicht nicht hundertprozentig ernst nehmen. Ich muss also nicht annehmen, dass die Maschine



plötzlich etwas fühlt, wenn ich sie so programmiere, dass sie eben so einen KI-Algorithmus ausführt oder dass sie da wirklich Bewusstseinsinhalt hat und Wahrnehmung und so weiter. Aber es macht zumindest Sinn, so zu tun als ob. Es gibt auch philosophisch gesehen vielleicht sogar ganz gute Argumente, warum man das nicht ganz ernst nehmen muss, aber es macht, wie gesagt, Sinn, so zu tun als ob. Und dann kann ich mir sinnvollerweise die Frage stellen, die Herr Krenn vorhin noch angesprochen hat: Was lernt denn dieses System eigentlich in dem Prozess, in dem es meine Aufgabe ausführt. Sagen wir mal, die Aufgabe ist eben genau wie bei AlphaFold. Ich kriege Aminosäureketten zusammen mit anderen Aminosäureketten, die evolutionär verwandt sind und soll daraus dann ableiten – anhand von einigen Vorbildern – wie das Protein hinterher aussehen wird, das sich daraus faltet. Dann ist es ja so, irgendwo lernt die Maschine, das gut auszuführen. Und man kann jetzt zum Beispiel zeigen, wenn man so gebastelte Strukturen nimmt, die es eigentlich gar nicht geben kann und man gibt dem System nur noch eine einzelne Aminosäurekette, dann findet man, dass hinterher bei der Struktur, die daraus gebildet wird, die Unsicherheit relativ groß ist, wenn es eine Struktur ist, die Quatsch ist und dass sie relativ gering ist, wenn es eine Struktur ist, die auch durchaus plausibel ist. Das heißt, es liegt irgendwo die Vermutung nahe – so ist es zumindest mal in einer Publikation geäußert worden –, dass ein System wie AlphaFold auch Informationen lernt, die über das hinausgehen, was ich vorne reingesteckt hab. Es lernt nicht nur die Coevolutions-Information zwischen den Aminosäureketten, sondern es lernt auch irgendwo vielleicht die physikalische Information, die mir sagt, welche Strukturen sind eigentlich möglich und welche nicht. Und die haben wir vorher gar nicht reingesteckt. Das ist nur eins von vielen Beispielen. Herr Krenn hat selber auch schon ein Beispiel benannt aus der Quantenoptik, in dem es ihm im Prinzip gelungen ist, auch neue Ansätze für experimentelle Aufbauten zu finden. Durch eine Analyse von einem KI-System, das eben diese Aufbauten, die Quantenzustände herstellen konnten, von denen man gedacht hatte, dass man die gar nicht machen kann, ausgespuckt hat. Und das ist eben das Spannende daran: Es ist jetzt nicht mehr nur intransparent, wie funktioniert denn das Ding im Ganzen? Das ist vielleicht bei komplizierten Simulationen auch irgendwann so. Wenn da viele mit dran schreiben, weiß man gar nicht mehr so ganz genau, wie funktioniert die Simulation eigentlich. Sondern hier ist auch noch zusätzlich intransparent: Was lernt denn das Ding eigentlich? Und das kann eben Information sein, die für uns relevant wäre, die wir als Wissenschaftler:innen gerne hätten. Und es gibt viele spannende Methoden, diese Informationen herauszufiltern. Zum Beispiel über weitere Algorithmen. Das nennt man manchmal symbolische Regression. Da kann man dann so eine Formel vielleicht extrahieren, die irgendwo das repräsentiert, was in dem System gespeichert wurde an Informationen, die ich noch nicht kenne. Und ein Problem entsteht sozusagen, wenn das jetzt aber Information ist, die ich noch gar nicht einordnen kann. Also vielleicht hat mein KI System sozusagen eine ganz neue Theorie entdeckt. Das ist eine Spekulation oder eine Hoffnung vielleicht auch, die manche Physiker tatsächlich auch schon in Publikationen geäußert haben. Und wie wäre ich dann als Mensch fähig, diese Theorie zu extrahieren? Wie wäre ich als Mensch fähig, diese Theorie herauszufinden mit meinen traditionellen Mitteln? Da kann es im Prinzip schon fundamentale Grenzen geben. Und das wäre natürlich irgendwo ein Verlust. Oder zumindest kann man sich auf den Standpunkt stellen, wenn man sagt: Na ja, wir wollen schon auch Anwendungen, wir wollen schon auch Sachen voraussagen können. Wir möchten natürlich gerne wissen, welche Aminosäurekette faltet zu welchem Protein, aber wir wollen auch gerne wissen, warum tut sie das? Wir wollen auch gerne wissen, was ist der Mechanismus dahinter? Wir wollen das wirklich verstehen. Und das ist auch ein Bruch, der uns besonders in der Philosophie interessiert. Kann es sein, dass die KI uns an einen Punkt führt, an dem die Voraussage, die Klassifikation, das genaue Erstellen von Korrelationen, die wir nutzen können, über das hinausgeht, was wir eigentlich verstehen können? Das wäre irgendwo ein epistemischer Verlust, würde man sagen, also ein erkenntnistmäßiger Verlust in der Wissenschaft, der gleichzeitig damit einhergeht. Vielleicht kann ich noch ganz kurz ergänzend sagen: Es ist sehr interessant, dass dieser Bericht von dem Joint Research Centre der Europäischen Kommission, den sie uns geschickt haben, auch sehr viele wissenschaftsphilosophische Grundfragen an einigen Stellen aufgreift. Man sieht hier einfach, dass die KI durch die ganzen



Folgen, die sie hat für den Forschungsprozess – auch für die Möglichkeit eines KI-Forschenden wie Herr Krenn auch schon angedeutet hat am Ende –, dass sie einfach viele dieser Fragen noch mal neu aufwirft und dass wir als Philosophinnen und Philosophen da noch mal ganz neu drüber nachdenken müssen: Wie funktioniert das denn eigentlich alles im Ganzen?

Moderatorin [00:20:36]

Mhm vielen Dank. Also einmal auf jeden Fall, dass KI die Möglichkeit schafft, Dinge in einem anderen Blickwinkel zu betrachten, den man vielleicht vorher nicht hatte. Aber andererseits auch, dass das ein Problem sein kann, wenn man dann eben nicht den Weg vorgezeichnet bekommt, sondern vielleicht nur das Ziel. Aber eigentlich die Intention ist ja in der Wissenschaft auch zu wissen, wie man zu einem Ergebnis gekommen ist. Wir haben schon diverse Fragen von außerhalb bekommen und zwar besteht anscheinend großes Interesse am praktischen Nutzen von dem KI-Einsatz. Und da ist einmal die Frage, gerne an Sie, Herr Krenn, inwiefern es schon Paper gibt, in denen KI zum Beispiel Experimente ersonnen hat, die schon konkrete Ergebnisse liefern, wie haltbar die sind. Weil ich glaube, dass es da schon relativ viele gibt, die das vorgestellt haben, aber nicht unbedingt immer so, dass sich das auch bestätigt hat.

Mario Krenn [00:21:32]

Ja es gibt zum Beispiel in der Quantenphysik mehrere Experimente, vielleicht zehn, fünfzehn Experimente, die tatsächlich grundlegend von KI-Algorithmik designed worden sind und dann im Labor gebaut wurden. Das erste, das aus meiner Forschung gekommen ist – ich habe den PhD gemacht im Labor von Anton Zeilinger, der 2022 den Nobelpreis für Physik bekommen hat und er hat an Quantenverschränkung gearbeitet. Und in dieser Richtung haben unsere ersten Computeralgorithmen auch Experimente gefunden – und zwar für ganz konkrete hochdimensionale, vielteilchen-verschränkte Quantenzustände. Das Interessante war, die Lösung haben wir 2014 gefunden und es hat bis 2018 gedauert, bis diese Lösungen im Labor gebaut wurden. Es ist genau der Blueprint gebaut worden. Es hat vier Jahre gedauert und dann hat das uns erlaubt, die Eigenschaften, die wir vier, fünf Jahre davor sehen wollten, dass wir die tatsächlich im Labor gesehen haben. Das ist ein Beispiel von vor einigen Jahren. Ein anderes Beispiel, das ist ganz neu, ist in einer Zusammenarbeit mit LIGO. Das ist ein internationales, großes Team. Die bauen einen Gravitationswellen-Detektor. Das sind Schwingungen in der Raumzeit, die von der Allgemeinen Relativitätstheorie vorhergesagt werden und zum Beispiel entstehen, wenn schwarze Löcher kollidieren. Und dort haben wir jetzt auch Lösungen gefunden, die anscheinend viel sensitiver sind als die menschlichen Lösungen. Und wir müssen jetzt nicht mehr unseren KI-Programmen vertrauen, die uns sagen, dass diese Lösungen sensitiver sind, weil wir jetzt die Lösung haben. Dann können wir diese Lösungen mit vielen anderen Methoden analysieren. Wir können die quasi per Hand berechnen und schauen, haben die tatsächlich die Eigenschaften? Und die scheinen die Eigenschaften zu haben, dass sie viel sensitiver sind, als was Menschen gemacht haben. Und hier vielleicht eine kleine Verbindung zu dem, was Professor Boge gesagt hat: Früher, also in der Quantenphysik, haben wir öfter aus den Lösungen direkt größere Konzepte ableiten können. Das heißt, wir haben dann viel mehr verstanden, haben per Hand viel mehr machen können. Für Gravitationswellen-Detektoren haben wir diese neuen Lösungen, die besser sind als das, was Menschen gefunden haben, analysiert für sechs Monate. Und wir haben nicht verstanden nach diesen sechs Monaten, warum die funktionieren, obwohl wir sie quasi per Hand berechnen können. Und das ist eine extrem frustrierende Situation für uns als Wissenschaftler. Auf der einen Seite wir haben Lösungen, die sind gut und wir wissen aber auch, dass diese Lösungen irgendwelche Ideen beinhalten müssen, quasi per Definition, die wir als Mensch nicht kennen. Und wir schaffen es nicht, diese Ideen zu extrahieren. Und wir sind da wirklich auf einem Weg in eine ganz neue Ära, in der wir aufpassen müssen, dass wir diese Grundmotivation der Wissenschaft, nämlich das Verstehen, dass wir das nicht verlieren. Und wir müssen ganz hart daran arbeiten, dass



wir das irgendwie noch reinbringen. Und das ist genau, warum die Arbeit von Professor Boge auch sehr, sehr wichtig ist.

Moderatorin [00:25:20]

Ich glaube direkt, Herr Boge, wollten Sie da was dazu sagen? Sie haben so aufgeregt mit dem Kopf genickt.

Florian Boge [00:25:23]

Ja, ich wollte das nur kurz ergänzen. Also es ist ganz vollkommen richtig – und das sehen wir auch in anderen Bereichen der Wissenschaft, wir arbeiten zum Beispiel mit Teilchenphysikern in Aachen so ein bisschen zusammen und da gibt es eine ganz ähnliche Beobachtung: Ein relativ einfaches und erstmal so überschaubares KI-System, so ein Entscheidungsbaum, der dann durch so bestimmte Methoden, effizienter gemacht wird. Da hat man festgestellt – das war dann auch auf Basis vorheriger Untersuchungen mit neuronalen Netzwerken – dass der mit bestimmten Beschreibungen eines Detektors super gut arbeiten kann, die physikalisch erstmal überhaupt keinen Sinn ergeben. Und da haben wir es genauso erlebt. Der Physiker, der uns das erzählt hat, hat gesagt: „Ach, das haben wir erst vor kurzem rausgefunden. Ich frage jetzt mal meine Postdoktorandin, ob sie schon eine Idee hat.“ Und sie haben bis heute nicht erklären können, was eigentlich der physikalische Prozess ist, der zugrunde liegt und der genau diese Beschreibung des Detektors so effizient macht. Also es gibt wirklich ganz strukturanaloge Beispiele, auch in anderen Wissenschaftszweigen. Man entdeckt da irgendwas Neues und merkt, das funktioniert super gut. Man hat keine Ahnung, wieso das so gut funktioniert und was die zugrunde liegenden, physikalischen Prozesse sind, was die zugrunde liegende Erklärung eigentlich ist.

Moderatorin [00:26:21]

Haben Sie da vielleicht gerade ein konkretes Beispiel noch aus einem anderen Wissenschaftsbereich parat, wo Sie es gerade schon angesprochen haben?

Florian Boge [00:26:27]

Ja, wie gesagt, ich kann da nur auf die Teilchenphysik rekurren. Das sind sogenannte Energieflusspolynome. Oder man kann eben auch bei AlphaFold – das habe ich ja vorhin auch schon angesprochen, das ist Frau Schoeders Bereich, in dem sie bestimmt viel mehr noch zu sagen kann – kann man sich auch fragen, warum funktioniert das eigentlich so gut? Es wurden da schon viele Ideen in das Training eingesteckt und auch in die Datenformatierung. Man hat eben den Eindruck, es geht schon ein bisschen über das hinaus, was man sozusagen hineingesteckt hat. Und da hat man ja circa fünfzig Jahre versucht, solche physikalischen Modelle zu finden oder biophysischen, mit sehr eingeschränktem Erfolg im Vergleich eben zu diesem KI-Erfolg. Aber da würde ich gerne an Frau Schoeder übergeben, falls sie was dazu sagen möchte.

Moderatorin [00:27:06]

Ja, sehr gerne. Frau Schoeder, entweder direkt dazu, ob Sie da was wissen oder ob Ihnen das auch aufgefallen ist, dass da Ergebnisse auftauchen, die man sich erst mal nicht erklären kann. Und dann auch noch mal mit Blick auf die Praxis, wenn wir schon bei AlphaFold sind, inwiefern denn da schon Arzneimittel draus entstanden sind beziehungsweise inwiefern es da schon Studien gibt, die auf Ergebnissen von AlphaFold basieren, man das vielleicht als Endverbraucher schon spürt.



Clara Schoeder [00:27:30]

Ja, das sind, glaube ich, zwei sehr gute Fragen. Zuerst zu der ersten, vielleicht im Anschluss direkt an Herrn Boge. Ich glaube, AlphaFold ist tatsächlich so ein ganz wunderbares Tool, weil die Datenbasis eben nicht unendlich groß ist, sondern eben recht überschaubar. Also die Proteine, auf denen das Ganze trainiert wird, das sind zweihunderttausend. Das ist immer noch eine überschaubare Nummer. Und wir haben eben traditionell schon ein relativ gutes Verständnis von dem Problem gehabt, weswegen wir sozusagen jetzt AlphaFold nehmen können und diese Black-Box, die es eigentlich ist, jetzt anfangen zu verändern. Also wir verändern den Input, wir schauen, wie verändert sich der Output. Eine Sache, die Leute zum Beispiel direkt erstaunt hat, ist, warum AlphaFold manchmal in der Lage ist, für ein Protein zwei verschiedene sogenannte Konformationen herauszugeben, also oft bewegen sich Proteine zum Beispiel von an nach aus. Und man kann AlphaFold modifizieren, dass es beides ausgibt – aber auch nicht immer. Und manchmal geht das eben nicht. Und das ist genau etwas, was wir jetzt sehr genau untersuchen, um zu verstehen, warum weiß es das manchmal? Und meistens weiß es das nicht. Und was müssen wir ihm noch zeigen, damit es das sozusagen auch erfassen kann? Und das ist manchmal sehr einzelfallabhängig und manchmal aber hat es ein größeres Muster erkannt, das wir vielleicht vorher so gar nicht so bewusst gemacht hatten oder ein Paradigma war, das uns bekannt war, aber so haben wir es noch nicht formuliert. Und das beflügelt ganz vieles, wo wir dann neue Dinge ausprobieren können oder zum Beispiel auch für Design verwenden können, weil dieses alles, was wir dann beschreiben können mit AlphaFold, können wir direkt wieder umsetzen in das Design. Und das Design ist genau der richtige Punkt für, was bringt uns das tatsächlich schon von akutem Nutzen? Also mit diesen Technologien selbst, da wird zurzeit ganz, ganz viel, vor allem in der Präklinik, also vor der klinischen Entwicklung, designt. Ob und inwieweit das jetzt Einfluss nehmen wird auf Medikamente, ist tatsächlich eine sehr gute Frage. Es gibt ein Beispiel für ein Computergestützt entwickeltes Protein, was bereits eine Marktzulassung erhalten hat. Das ist ein Impfstoff. Der wurde allerdings mit klassischen computergestützten Methoden entwickelt. Also da wurde das Protein biophysikalisch designt und es handelt sich um einen Nanopartikel, der den Impfstoff trägt. Das Produkt heißt Scalcovion und ist in Südkorea zugelassen. Das ist so ein erstes Beispiel, bei dem das jetzt geklappt hat. Es gibt aber auch große Fragen, zum Beispiel: Wahrscheinlich ist es nicht so schlimm, wenn ich nur wenige Veränderungen am Protein einführe. Das machen wir auch traditionell oft. Aber was ist mit einem Protein, das der Computer komplett neu designt hat? Welche Eigenschaften hat es überhaupt und wie reagiert der Körper da drauf? Das ist eine Frage, der zurzeit viele versuchen nachzugehen. Ein Stichwort dort ist immer Immunogenität. Das ist auch bei normalen Wirkstoffen schon ein Problem. Wird es bei solchen Impfstoffen, oder bei solchen Wirkstoffen auch ein Problem sein? Das versuchen wir jetzt herauszufinden. Und da werden diese Wirkstoffe ganz normal den Weg auch der klinischen Zulassung beschreiten müssen.

Moderatorin [00:30:31]

Mhm. Alles klar, vielen Dank. Also ich höre auch da raus, es entstehen neue Dinge, bei denen man sich gar nicht unbedingt in jedem Fall so sicher ist, wo die herkommen. Und da habe ich auch noch eine etwas konkretere Frage zu, und zwar sind ja Trainingsdaten immer ein Abbild von dem, was schon bekannt war. Da ist es so intuitiv erst mal der Gedanke, dass daraus nicht unbedingt etwas Neues entstehen kann, aber es scheint ja doch so zu sein. Vielleicht, Herr Krenn, können Sie mit Blick auf das Große Ganze da noch mal was zu sagen?

Mario Krenn [00:31:04]

Ja, für gewisse Designfragen oder Entdeckungsfragen braucht man gar keine Trainingsdaten. Zum Beispiel in der Physik für das Design von Experimenten, das Finden von neuartigen Mikroskop-



Technologien oder Teleskop-Technologien benutzt man gar keine Trainingsdaten. Was wir stattdessen benutzen, sind physikalische Simulatoren, die quasi die Grundgleichungen der Physik, wie zum Beispiel die Gleichung von Schrödinger, Gleichung von Maxwell, Einsteins Gleichungen, die die simulieren können. Das heißt, wir haben ein Programm, dem geben wir ein hypothetisches Experiment und dieses Programm kann ganz genau sagen, basierend auf allem, was wir jetzt aus der Natur wissen, von den Grundgleichungen, wie sich dieses Experiment verhalten wird. Und wenn man so einen Simulator hat, kann man Millionen von Experimenten diesem Programm geben. Und dieser Simulator zeigt uns dann, welches dieser Million Experimente die höchste Sensitivität hat oder beste Auflösung erreichen kann. Das ist vollkommen ohne Trainingsdaten. Das funktioniert in der Physik gut, das funktioniert teilweise in der Chemie. Da kann man die Gleichungen auch noch berechnen, zum Teil bei gewissen Fragen. Das heißt, da kann man auch datenfreie Exploration und Computer-gestützte Entdeckungsreisen machen. Ich glaube, das wird schwieriger, wenn man zu komplexen großen chemischen Systemen geht. Und ich glaube, in der Biologie endet das dann. Das ist vollkommen unmöglich, dass man dort von den Grundgleichungen der Physik wirklich die Eigenschaften berechnet, weil die Systeme zu komplex sind.

Moderatorin [00:33:00]

Mhm. Ich weiß nicht, wahrscheinlich kann man das nicht mit einer kurzen Antwort wiedergeben, aber vielleicht Herr Boge, aus wissenschaftsphilosophischer Sicht: Kann KI so was wie kreativ sein?

Florian Boge [00:33:16]

Na ja, in gewisser Weise, sehen Sie es mal so: Sie haben gesagt, die Daten, da würden wir annehmen... also bleiben wir jetzt mal bei KI-Systemen, die wirklich auf echten Daten trainiert werden. Man kann die auch simulationsbasiert mit Daten trainieren, aber nehmen wir mal an, es wird auf echten Daten trainiert, die wir wirklich empirisch erhoben haben. Die Frage ist ja sozusagen: Funktioniert die KI ähnlich wie ein Mensch? Das heißt, sieht die in Anführungszeichen das Gleiche in den Daten wie ich? Es gibt da vielleicht eine ganz interessante, ganz nette Metapher von Geoffrey Hinton in Reaktion auf den Erfolg von diesen Large Language Models, die also flüssig Englisch oder eben auch Deutsch sprechen können – zumindest schriftlich oder in einigen Fällen eben auch Audio-Dateien ausgeben können. Er hat gesagt, es ist, als wären Aliens auf der Erde gelandet, nur keinen interessiert's, weil sie einfach flüssig Englisch sprechen. Man kann es im Prinzip jetzt schon so sehen, dass die künstliche Intelligenz insofern, als sie eben die Intelligenz mimt, einfach eine ganz andere Art von Intelligenz ist, ja? Oder eine ganz andere Art ist, sozusagen darauf zu gucken. Wenn Sie sich jetzt einfach zwei menschliche Betrachter anschauen, die auf die gleichen Sachen gucken, dann werden die auch ganz unterschiedliche Sachen feststellen. Das heißt, wir haben jetzt sozusagen noch mal eine ganz andere Drittinstantz, die darauf gucken kann und die dementsprechend auch ganz andere Dinge darüber hinaus finden kann. Ich glaube, das ist ein Punkt, wo man, ich weiß nicht, ob man's Kreativität nennen muss, aber wo man sozusagen noch mal einen zusätzlichen Blick bekommt und dadurch vielleicht auch irgendwo etwas Neues relativ zu dem, was wir bereits wissen.

Moderatorin [00:34:33]

Mhm. Alles klar. Also könnte man doch eine relativ kurze Antwort darauf geben. Da war ich gar nicht drauf vorbereitet. Noch eine Frage zum Ziel der Wissenschaft. Dadurch, dass KI ja jetzt gegebenenfalls da auch beim Erkenntnisprozess vielleicht Sachen wegnimmt, merkt man, dass das Ziel sich ändert innerhalb der Wissenschaft? Weg vom „Ich möchte verstehen“, hin zu Anwendungen, die vielleicht relevanter sind? Vielleicht erst mal Herr Krenn?



Mario Krenn [00:35:11]

Ja, also, wir haben jetzt Technologien und KI-Systeme, die uns helfen, neue Technologien zu entwickeln. Die Physik möchte quasi das Universum und alles, was es beinhaltet, verstehen und beschreiben können. Das heißt, wenn wir jetzt neue Lösungen haben, dann möchten wir diese natürlich auch verstehen. Und das ist jetzt genau ein Problem. Wir wissen noch nicht genau, wie wir diese Lösungen von dieser ganz anderen Intelligenz, wie wir diese jetzt verstehen können. Das wäre fast wie, wenn ein Alien kommt und gibt uns eine Lösung und wir müssen das dann verstehen. Es ist uns im Moment noch nicht klar, wie wir diesen Prozess machen: Von, wir haben eine Lösung, bei der wir sehen, dass sie funktioniert, dass wir das zurückspringen auf das Verständnis. Also es beschäftigen sich viele Leute damit. Das heißt, Verständnis selbst ist immer noch ganz oben auf der Karte der Wissenschaftler. Es ist nur im Moment nicht ganz klar wie wir diese neuen Systeme analysieren oder verstehen können. Potenziell, das hoffe ich, könnte man selbst KI-Systeme entwickeln, die quasi wie ein Lehrer für uns wirken, die diese Systeme verstehen und dann runterbrechen können auf Grundkonzepte, die wir als Menschen dann wieder verstehen können. So etwas gibt's im Moment noch nicht, aber ich glaube, so etwas wird in der relativ nahen Zukunft wirklich wichtig werden.

Moderatorin [00:37:01]

Okay, danke schön. Eine Frage vielleicht an Sie, Frau Schoeder, aus der Praxis. Dann einer der Herren, die vielleicht noch mal größer was dazu sagen können, also mit einem breiteren Überblick: Werden alle Forschenden in Zukunft mehr dazu ausgebildet werden müssen, den KI-Prozess nachzuvollziehen? Vielleicht merken Sie das bei Ihnen auch schon im Labor, anstatt selbst Experimente zu entwickeln und Daten zu erheben?

Clara Schoeder [00:37:28]

Ja, das ist eine sehr gute Frage, wie das auch in der Ausbildung von unseren Wissenschaftlern und Wissenschaftlerinnen eine Rolle spielen wird. Was wir schon tatsächlich merken, ist, dass oft sozusagen dieser allererste Schritt, zum Beispiel um Skripten oder Coden zu lernen, der ist sehr viel einfacher geworden, weil ich kann einfach ein Large Language Modell fragen, Teile des Codes für mich zu erstellen. Das führt natürlich aber dazu, dass ich oft nicht unbedingt verstanden habe, wie das Ganze funktioniert. Für uns ist es tatsächlich umso wichtiger, die Grundlagen wirklich zu vermitteln und sicherzustellen, dass die Anwenderinnen und Anwender wirklich auch ein Verständnis davon entwickelt haben. Wir haben ein ganz ähnliches Problem beim Arbeiten im Labor oft tatsächlich, dass wir heutzutage vieles automatisiert haben. Das heißt, ich stelle irgendwie vorne was rein und hinten kommt irgendwas raus und das verwerte ich dann weiter. Aber was dazwischen passiert, hat oft einen ganz, ganz großen Einfluss auf den Prozess und das muss ich verstanden haben. In der Ausbildung versuche ich vor allem, wirklich diese Grundlagen wieder aktiver zu leben, einfach damit man besser trouble-shooten kann. Was dann passiert in zum Beispiel dem Skript, das ich mit ChatGPT geschrieben habe oder bei der Anwendung im Labor, wo ich mich vielleicht darauf verlassen habe, dass die Auswertesoftware das richtig gemacht hat, muss ich trotzdem verstehen, wo der Fehler sein kann. Und das ist tatsächlich auf jeden Fall eine von den ganz großen Herausforderungen an uns in der Lehre und wir müssen da sicherlich viel auch an uns auf unserer Seite als Lehrender arbeiten, um das zu integrieren, dass wir wirklich da auch wachsam bleiben. Ich wette, die Kollegen haben da bestimmt auch noch ein paar Ideen zu.

Moderatorin [00:38:56]

Ja, haben die Kollegen?



Florian Boge [00:38:58]

Also ein Punkt, den ich, total unterstreichen kann, den Frau Schoeder gerade gemacht hat, ist natürlich, dass man in der Lehre auch eine gewisse KI-Kompetenz irgendwo erlernen muss oder dass man die vermitteln muss, dass das ganz, ganz wichtig ist. Die Forschenden, die schon sehr viel Erfahrung haben, die sind sich meistens der Probleme auch sehr bewusst. Es gibt ja nicht nur die Möglichkeit, dass die KI was Neues entdeckt und uns deswegen so tolle Voraussagen gibt, sondern es gibt ja auch die Möglichkeit, dass sie, na ja, zum Beispiel einen Hack entdeckt, sozusagen. Also die Daten enthalten irgendwo einen kleinen Hinweis darauf, welches Protein, sagen wir jetzt mal wieder, um bei dem Beispiel zu bleiben, welche Struktur haben muss. Und das hat eigentlich gar nichts mit dem echten Faltungsmechanismus oder irgendwelchen evolutionären Verbindungen zu tun. Ein Beispiel ist, es gab das mal ganz früh bei der KI, dass Bilder von russischen Panzern immer verpixelt waren und dann wurden alle verpixelten Bilder von Panzern als russische Panzer klassifiziert. Es gibt also eine Korrelation, ohne dass es einen Kausalzusammenhang gibt. Und das muss man ja auch wissen. Man nennt das manchmal das Clever-Hans-Verhalten, weil es da mal so ein Pferd gab, das schien irgendwie rechnen oder zählen zu können, hat in Wirklichkeit aber nur geguckt, ob der Trainer ganz subtile Signale zeigt: „Ach, jetzt ist es aber auch gut. Jetzt kannst du aufhören zu klopfen mit deinem Fuß.“ Und dieses Problem gibt's eben auch, neben dem Problem, dass wir vielleicht an Verstehensgrenzen stoßen, weil die KI mehr kann, gibt es auch das Problem, dass sie vielleicht einfach viel weniger macht, als wir denken, dass sie tut. Und das ist zum Beispiel den Studierenden oft gar nicht klar. Wir haben jetzt in der Philosophie ganz oft das Problem, dass Studierende ihre Hausarbeiten und sogar Praktikumsberichte von KI schreiben lassen und wir entdecken dann irgendwann, da kann was nicht stimmen, weil mittelalterliche Philosophen sich über physikalische Theorien des 20. Jahrhunderts auslassen in den Arbeiten. Also das kann einfach zeitlich nicht hinkommen. Oder, Referenzen werden falsch zitiert und man wundert sich dann: „Na, das hat sie doch im Referat richtig gemacht“, oder er. Und, am Ende kommt da absolut grober Unfug raus. Und ich glaube, diese KI-Kompetenz zu trainieren, dass auch gerade Studierende oder junge Forschende lernen: Was sind die Möglichkeiten? Was sind aber auch gleichzeitig die Grenzen? Worauf muss ich achten und wovor muss ich mich in acht nehmen? Wie wichtig bleibt weiterhin empirische Überprüfungen oder Überprüfungen anhand von bekanntem theoretischem Wissen? Das ist, glaube ich, ganz, ganz entscheidend und das habe ich auch immer wieder gehört im Austausch mit Wissenschaftler:innen, dass sie beobachten, dass es in manchen Fällen nicht gelingt, dass manche Forschende einfach zu naiv das glauben, was die KI ausspuckt und das zu wenig hinterfragen. Während, wie gesagt, in der Regel erfahrenere Forschende da einen etwas kritischeren Blick haben und das ist auch gut so, denke ich.

Moderatorin [00:41:20]

Hm-hm. Umso schwieriger können ja zum Beispiel dann mit diesem Blickwinkel auch vollautomatisierte Labore sein. Das hatten Sie am Anfang, glaube ich, schon mal ganz kurz angesprochen, Herr Krenn. Inwiefern werden denn solche vollautomatisierten Labore schon in der Forschung eingesetzt, beziehungsweise was ist da in Zukunft absehbar und erstrebenswert?

Mario Krenn [00:41:40]

In der Chemie gibt es sehr viele Arbeiten zu vollautomatisierten Laboren. Wahrscheinlich gibt es das in der Biologie auch. Da kann wahrscheinlich Professor Schoeder mehr dazu sagen. Weil in der Chemie ist es so, dass es sehr, sehr schwierig ist, die Grundgleichungen zu berechnen. Das heißt, man kann seinen Resultaten von Simulatoren oft nicht vertrauen. Und wenn man jetzt eine Exploration startet und neue Materialien suchen möchte für Solarzellen zum Beispiel oder für Batterien, und den Resultaten aber nicht wirklich vertrauen kann, dann kann man der ganzen



Exploration nicht wirklich trauen. Deswegen werden jetzt sehr viele dieser In-Silico-Simulationen ersetzt durch wirkliche Roboter, die dann eine wirkliche Synthese im Labor machen und die Eigenschaften tatsächlich messen, weil die Natur gibt am Schluss die richtige Antwort. Das heißt, in der Chemie gibt es hier sehr viele Forschungsgruppen. Es gibt auch in Deutschland eine berühmte Forschungsgruppe in Erlangen von Christoph Brabec. In Kanada, am MIT gibt es viele Gruppen. Es gibt auch einige Start-ups, die sich jetzt gegründet haben. Genau, weil die sehen, dass durch Materialforschung es gewaltige Erkenntnisgewinne geben könnte, die sich direkt praktisch umsetzen lassen. Und die haben auch gesehen, dass die Technologie heute so weit ist, dass man das wahrscheinlich praktisch in den nächsten paar Jahren tatsächlich machen kann.

Moderatorin [00:43:35]

Das heißt, es ist keine so große Zukunftsmusik mehr, sondern es gibt Labore, die das schon aktiv einsetzen beziehungsweise kurz davor sind, damit Resultate zu erzielen?

Mario Krenn [00:43:44]

Ja, es gibt Resultate, die auch schon von automatisierten Laboren gemacht worden. Es gibt, glaube ich, noch keine wirklich fundamentale Breakthroughs. Ich bin in Kanada Postdoc gewesen und im Nachbarraum war ein Labor und dort haben Roboter Trainingsdaten für sich selbst aufgenommen, die ganze Nacht. Also das, das läuft schon seit einigen Jahren.

Moderatorin [00:44:18]

Mhm. Da hab ich auch noch eine Frage an Frau Schoeder dazu. Und zwar ist es ja zum Beispiel durch AlphaFold oder durch Proteindesign-Methoden einfacher möglich, herauszufinden welche Proteine für was zuständig sind. Ist das der Knackpunkt in der Debatte? Oder ist jetzt die Schwierigkeit vielleicht noch die Herstellung von den Proteinen, weil sich da ja auch Leute über die Sicherheit Gedanken machen? Also reicht die KI, dann weiß ich, wie's geht und dann ist es einfach, oder brauche ich noch ein kompliziertes Labor im Hintergrund, um Viren oder chemische Waffen herzustellen?

Clara Schoeder [00:44:52]

Ja, das sind, glaub ich, ganz viele Fragen auf einmal. Tatsächlich ist, glaube ich, das für uns eine Erkenntnis. Jetzt, wo wir Struktur vorhersagen können, stellen wir halt fest, aus der Struktur können wir nur so viel ableiten und das nächste ist tatsächlich zu verstehen: Was ist denn die Funktion? Weil wenn ich ein unbekanntes Protein habe, oft kann ich gar nicht besonders viel über die Funktion sagen. Wir haben das tatsächlich auch selbst erlebt: Wir hatten bestimmte Proteine gesucht und dann haben wir's im Labor untersucht und dann war es so: Ja, was machen die denn überhaupt? Aber das ist zum Teil heutzutage andersrum gar nicht einfach festzustellen. Und zu der anderen Frage ist auch, wenn wir jetzt Proteine, zum Beispiel, de-novo designen, ist tatsächlich auch oft die Frage: Designen wir vielleicht etwas, was wir gar nicht haben designen wollen, weil es vielleicht auch in diesen Algorithmen mit drin war? Und das ist, glaube ich, so überhaupt noch nicht abzusehen. Es gibt halt, glaub ich, große Befürchtungen. Also es ist heute auch immer noch schwer, ein Protein zu designen und herzustellen, vor allem auch in größeren Maßstäben. Und das heißt, diesen Proteinherstellungsprozess, den werde ich nicht los und den werde ich auch immer noch weiter brauchen. Und das ist am Ende des Tages auch Biologie, die auch ein bisschen Input aus dem Labor wirklich braucht, also auch laborspezifisches Wissen braucht. Das heißt, diese Barriere, um Dinge herzustellen, die ist auf jeden Fall auch immer noch da. Ich werde tatsächlich relativ viel auch nach den Risiken gefragt und auf der einen Seite glaube ich, dass das noch nicht so richtig



abzusehen ist. Auf der anderen Seite, ich vergleiche's auch immer ganz gerne mit der Chemie: Auch in der Chemie muss man sich jeden Tag entscheiden, was es denn ist, was man synthetisiert. Und ich glaube, da kommen wir einfach im Proteindesign auch hin. Wir müssen auch jeden Tag eine Entscheidung treffen, was das ist, worauf wir arbeiten und wofür wir Dinge anwenden und wie wir auch überprüfen, wie sie zum Beispiel in die Umwelt gelangen und wie wir auch testen, welche Eigenschaften sie haben. Und ich glaube, das ist sehr analog tatsächlich zu dem, was die Chemie vor vielen Jahren schon hatte und immer oder jeden Tag auch diskutieren muss. Insofern, also, ich glaube, wir werden auf der einen Seite viel lernen, was in den Algorithmen ist, natürlich. Auf der anderen Seite würde ich es aber auch nicht überschätzen. Also auch heute noch ist es schwierig, ein Protein zu designen. Es ist nicht so, man drückt nur auf den Knopf und dann kommt ein perfektes Protein raus. Das braucht auch immer noch viel Testung, viel Charakterisierung, viel Optimierung.

Moderatorin [00:47:17]

Mhm. Angenommen, die Forschung mit oder ohne KI wäre irgendwann so weit, dass aus der Faltung oder aus der Aminosäuresequenz klar ist, welche Eigenschaften ein Protein hat. Also wenn ich eine Eigenschaft möchte und ich kann ein Protein erstellen, wäre es dann trotzdem noch komplex, das wirklich herzustellen?

Clara Schoeder [00:47:38]

Vor allem, wenn ich die Eigenschaften sehr gut einschätzen kann. Zum Beispiel welches Expressionssystem, posttranslationale Modifikation, ist auch ein Stichwort immer, das sehr wichtig ist und wie ich's aufreinigen kann. Proteinherstellung ist heutzutage... Es ist eine wichtige Technologie, aber sie ist auch durchaus etwas, was, im Labor gut durchführbar ist. Es gibt Proteine, die sind schwieriger herzustellen als andere. Es ist relativ einfach, ein globuläres, lösliches Protein einer geringen Aminosäureanzahl herzustellen. Je größer, desto komplizierter. Je komplexer das System, desto komplizierter, wie zum Beispiel Membranproteine. Es ist immer noch schwieriger, es ist sehr, sehr schwierig, tatsächlich sogar. Und, also es kommt aber so ein bisschen drauf an, welches System ich bearbeite und wie viel Wissen ich auch darüber habe.

Moderatorin [00:48:25]

Mhm. Okay, also trotzdem noch komplex, aber eine große Hürde ist überhaupt herauszufinden, welche Eigenschaften ein Protein mitbringen soll. Wir haben noch eine Frage an Herr Boge. Mit Blick auf die Zeit: Ich hatte mit Ihnen allen schon abgeklärt, dass es in Ordnung ist, wenn wir zehn Minuten überziehen. Aber nur, dass Sie Bescheid wissen, wir haben jetzt in etwa 13 Uhr gleich. Das heißt, wenn doch noch irgendwas ist, sagen Sie Bescheid. Trotzdem die Frage an Herrn Boge, von den Journalist:innen auch noch: Sind wir gerade dabei, unter dem Dach der Wissenschaft unsere Unmündigkeit zu verschulden, wenn wir die KI über unser eigenes Verstehen, über das Erklärbare hinaus, gewähren lassen?

Florian Boge [00:49:06]

Natürlich eine sehr, sehr philosophische Frage. Ich kann sie eigentlich nur sehr unphilosophisch beantworten. Ich bin ja Wissenschaftsphilosoph und setze mich mit diesen ganz elementaren Fragen ganz wenig auseinander. Also, ich würde sagen, ich weiß nicht, ob wir von Unmündigkeit reden müssen. Im Grunde genommen, aktuell haben wir es ja immer alles noch weitestgehend im Griff. Es gibt natürlich schon ein bisschen die Gefahr der Verselbstständigung. Ich hab da kürzlich mit jemandem drüber gesprochen, der aus der Physik kommt und auch viele interessante



Algorithmen entwickelt hat. Dass er grundsätzlich schon die Gefahr sieht, dass zum Beispiel, weiß ich nicht was, solche Horrorszenarien, dass man verschiedene KI-Systeme vernetzt, was grundsätzlich möglich ist, und dass die dann irgendwie Zugriff auf automatisierte Waffen wie eben Drohnen bekommen. Das ist natürlich denkbar. Diese Gefahren sollten wir auf jeden Fall immer irgendwo im Blick behalten. Wenn wir jetzt rein bei der Wissenschaft bleiben, dann sehe ich die Gefahr erstmal noch nicht so, denn wir haben ja gerade schon auch von Professor Krenn und Professor Schoeder gehört, dass an vielen Stellen der Mensch sozusagen immer noch einschreitet und einschreiten möchte und auch erstmal schaut, was ist denn das, was ich da jetzt eigentlich rausgefunden habe? Insofern sehe ich das noch nicht so ganz. Also wir haben sozusagen nur die Gefahr, dass wir sehr lange brauchen werden, dass wir eigentlich in gewisser Weise totale Fortschritte machen, praktische Fortschritte, bestimmte Probleme ganz einfach lösen können, dass wir aber an anderer Front total zurückbleiben und vielleicht auch an der Front, die eben die Wissenschaft zur Wissenschaft macht, nämlich dass wir eben das wissenschaftliche Verstehen einfach zurücklassen. Also ich würde da weniger von Unmündigkeit reden, als mehr von einem epistemischen oder erkenntnistmäßigen Verlust in gewisser Hinsicht. Manche Sachen können wir jetzt ganz wahnsinnig schnell, andere Sachen können unfassbar lange brauchen. Und das ist sozusagen ein bisschen der Konflikt, dass man sich so einseitig weiterentwickelt, dass man dann einseitig sehr weit voranschreitet und auf anderer Seite, wo man das eigentlich gerne gehabt hätte, ziemlich lange warten muss oder vielleicht auch nie weiterkommt.

Moderatorin [00:50:50]

Das ist ja mit Sicherheit auch was, was die Wissenschaftler:innen selbst beschäftigt. Also was nicht nur in der Wissenschaftsphilosophie angesprochen wird, oder? Bekommen Sie da auch was von mit?

Florian Boge [00:50:57]

Ja, selbstverständlich. Also ich meine, im Grunde genommen können Sie alles, was Professor Krenn gerade erzählt hat, ja genau darauf ummünzen. Dass er sowohl den Fall gesehen hat. In der Pharmakologie haben wir das zum Beispiel auch, dass man mithilfe von KI wirklich neue Sachen entdecken kann, die man dann auch empirisch überprüfen kann. Es gibt eben zum Beispiel Pharmazeutika, die machen zwei verschiedene Sachen, also die docken an zwei verschiedene Rezeptoren im Körper an oder nur an einen. Und wenn man die jetzt klassifizieren lässt von einer KI, dann kann das ziemlich akkurat funktionieren, dass die Klassifikation klappt. Und dann kann man mit solchen Zusatzmethoden, die nennen sich erklärbare KI-Methoden, dann schauen, was hat die KI da eigentlich auf atomarer Ebene sich genau angeschaut? Und dann findet man zum Beispiel, aha, die Chemikalie hat Koffein als Substruktur und sieht dann genau, das ist sozusagen der entscheidende Prädiktor. Empirisch ist das auch schon mal so rausgefunden worden, dass Koffein tatsächlich als Substruktur erzeugt, dass eine Chemikalie zwei verschiedene Sachen macht, an zwei verschiedene Arten von Rezeptoren andockt. Das machen zum Beispiel die Lebenswissenschaftler in Bonn oder eben Herrn Krenns eigenes Beispiel, dass er es geschafft hat, Quantenzustände herzustellen, von denen er ursprünglich dachte, die kann es gar nicht geben. Weil ein ganz neuer Ansatz verwendet wurde, sozusagen, um experimentelle Aufbauten oder Teile von experimentellen Aufbauten zu kombinieren. Den Ansatz kannte man schon, aber die KI hat ihn nicht erzählt bekommen und hat ihn sozusagen wiederentdeckt und ganz neu verwendet. Das heißt, das haben wir auf der einen Seite. Wir haben eben diese, diese Beispiele, wo es diese, diese eklatanten oder beeindruckenden Fortschritte tatsächlich auch im Verstehen gibt, ja? Wir lernen was ganz Neues und wir können dadurch auch mehr verstehen darüber, wie die Welt funktioniert. Auf der anderen Seite aber eben auch die Beispiele Gravitationswellen, Astronomie wurde gerade benannt oder auch Teilchenphysik an bestimmten Stellen und eben auch Proteinstrukturbiologie, wo wir nur wissen, es funktioniert irgendwie ganz besonders toll und es muss da irgendwie was



Interessantes geben. Aber was ist das denn eigentlich? Vielleicht noch eine ganz kleine Anekdote, die ich dazu erzählen kann. Ich habe mit einem Physiker kürzlich gesprochen, der solche Algorithmen entwickelt, um so Formeln, die ein Mensch verstehen könnte, zu extrahieren aus einem KI-System, zum Beispiel aus einem neuronalen Netz. Also der versucht auszulesen mit einem bestimmten Algorithmus, was denkt sozusagen das Netz da und wie kann ich das durch eine Formel beschreiben? Und er macht das bereits seit acht Jahren, aber ist in den Anfangstagen von seinen Physikkollegen und -vorgesetzten dafür vollkommen verlacht worden, weil man das gar nicht ernst genommen hat. Und ich würde sagen, na ja, eigentlich lacht er jetzt sozusagen zuletzt. Denn das ist genau das, was gerade, ja, die Frontier Science sozusagen ist, zu verstehen: Was lernen diese KI-Systeme? Wie denken die in Anführungszeichen über die Welt und was können wir daraus lernen?

Moderatorin [00:53:27]

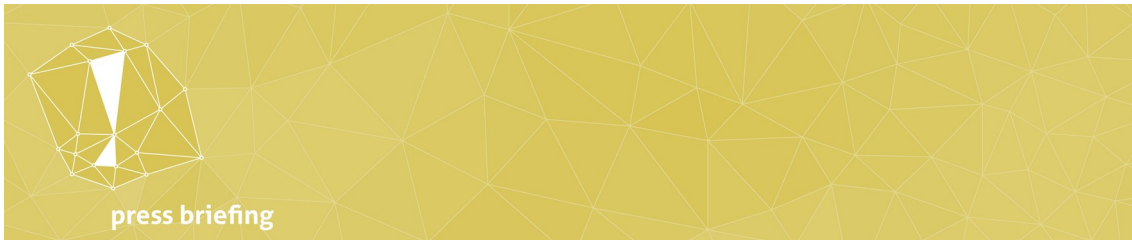
Hm, vielen Dank. Zeit für eine letzte Frage haben wir vielleicht noch. Ich habe zwar eigentlich angekündigt, dass es heute nur um den Forschungsprozess gehen soll und nicht um den Publikationsprozess, aber wir haben noch eine Frage zu der Konferenz reinbekommen von Stanford, die ich angekündigt habe. Und zwar gerne an Sie vielleicht, Herr Krenn. Bei der morgigen Open Conference of AI Agents übernehmen KI-Systeme auch den Review-Prozess. Welche Risiken und Chancen sehen Sie darin, falls man das kurz beantworten kann?

Mario Krenn [00:53:58]

Ja, ich bin begeistert von KI-Anwendung in allen Bereichen der Wissenschaft, außer im Review-Bereich. Im normalen Review-Prozess, hat man quasi Menschen, die unterschiedliche Biases haben, aber da immer mehrere Reviewer zusammenkommen, sind diese Biases ziemlich divers und so weiter. Das heißt, da hat man nicht so ein großes Problem, dass man extreme Biases hat. Bei diesen automatisierten Reviewern hätte ich die Sorge, dass die extreme Biases haben und dass es auch möglich wird, wenn man ein Programm hat, das zwischen akzeptiert und nicht akzeptiert entscheiden kann. Dann wird es eine gewisse Wahrscheinlichkeit haben, achtzehn Prozent akzeptiert. Und dann könnte man das Programm nehmen, um einen Artikel zu optimieren, damit, nur indem man Wörter verändert, die Akzeptanzrate höher wird. Das ist ein extremer Bias. Ohne, dass man irgendwelche inhaltlichen Erkenntnisse dazu gibt, könnte man rein durch Manipulation von der Sprache die Akzeptanzrate nach oben treiben. Das ist ein ganz einfacher Bias. Es gibt viele, viele andere. Es kann sein, dass neue Ideen einfach immer rejected werden, weil sie zu neu sind oder weil sie zu unklar sind, obwohl diese neuen Ideen in der Zukunft große Auswirkungen haben können. Ich glaube, das gibt es beim Menschen auch, aber Menschen sind, die Community an Wissenschaftlern ist so divers, dass irgendwann die wichtigen Artikel auch publiziert werden, auch akzeptiert werden. Aber bei KI-Systemen habe ich da wirklich große Sorgen.

Moderatorin [00:56:01]

Okay, alles klar. Ja, perfekt. Das war jetzt schon ein Teaser auf ein vielleicht zukünftiges Science Media Center Angebot zum Publikationsprozess. Aber für heute sind wir am Ende. Deswegen vielen Dank an Sie, liebe Expertinnen und Experten und an Sie, liebe Journalistinnen und Journalisten, für Ihre Aufmerksamkeit und für Ihre Fragen. Im Laufe des Tages werden wir eine Audioaufzeichnung, eine Videodatei und ein maschinell erstelltes Transkript zur Verfügung stellen, vermutlich innerhalb der nächsten halben Stunde. Später am Nachmittag gibt es dann auch noch ein redigiertes Transkript zu diesem Meeting. Die Links dazu finden Sie in der Reminder-Mail von heute morgen, da können Sie dann darauf zugreifen. Ich bedanke mich bei Ihnen. Auf Wiedersehen, noch einen schönen Tag.



Florian Boge [00:56:43]

Tschüss und vielen Dank fürs Gespräch.



press briefing

Ansprechpartnerin in der Redaktion

Samantha Hofmann

Redakteurin für Digitales und Technologie

Telefon +49 221 8888 25-0

E-Mail redaktion@sciencemediacenter.de

Impressum

Die Science Media Center Germany gGmbH (SMC) liefert Journalisten schnellen Zugang zu Stellungnahmen und Bewertungen von Experten aus der Wissenschaft – vor allem dann, wenn neuartige, ambivalente oder umstrittene Erkenntnisse aus der Wissenschaft Schlagzeilen machen oder wissenschaftliches Wissen helfen kann, aktuelle Ereignisse einzuordnen. Die Gründung geht auf eine Initiative der Wissenschafts-Pressekonferenz e.V. zurück und wurde möglich durch eine Förderzusage der Klaus Tschira Stiftung gGmbH.

Nähere Informationen: www.sciencemediacenter.de

Diensteanbieter im Sinne MStV/TMG

Science Media Center Germany gGmbH

Schloss-Wolfsbrunnenweg 33

69118 Heidelberg

Amtsgericht Mannheim

HRB 335493

Redaktionssitz

Science Media Center Germany gGmbH

Rosenstr. 42–44

50678 Köln

Vertretungsberechtigter Geschäftsführer

Volker Stollorz

Verantwortlich für das redaktionelle Angebot (Webmaster) im Sinne des §18 Abs.2 MStV

Volker Stollorz



science
media center
germany