

Künstliche Intelligenz: unerwartete Ergebnisse

Forschende der Uni Bonn blicken hinter die Kulissen des Maschinellen Lernens in der Arzneimittelforschung

Künstliche Intelligenz (KI) ist immer weiter auf dem Vormarsch. Bislang erschienen die Anwendungen häufig wie eine Black Box: Wie die KI zu ihren Ergebnissen kommt, blieb oft im Verborgenen. Nun kommt Licht ins Dunkle: Prof. Dr. Jürgen Bajorath, Chemieinformatiker an der Universität Bonn, hat mit seinem Team eine Methode entwickelt, die enthüllt, wie bestimmte KI-Anwendungen in der Arzneimittelforschung vorgehen. Die Ergebnisse sind unerwartet: Die KI-Programme erinnerten sich weitgehend an ihnen bekannte Daten und gingen in der Vorhersage von Arzneimittelwirksamkeiten kaum auf den Kern spezifischer chemischer Wechselwirkungen ein. Die Ergebnisse sind nun in "Nature Machine Intelligence" publiziert.

Welches Arzneimittel-Molekül hat die beste Wirksamkeit? Forschende suchen fieberhaft nach effizienten Wirkstoffen, um Krankheiten zu bekämpfen. Häufig docken diese Präparate an ein Protein an. Dabei handelt es sich zumeist um Enzyme oder Rezeptoren, durch die eine bestimmte Wirkungskette ausgelöst wird. Teilweise sollen bestimmte Moleküle auch unerwünschte Reaktionen im Körper blockieren – etwa eine überschießende Entzündungsreaktion. Angesichts der Fülle verfügbarer chemischer Verbindungen gleicht diese Forschung auf den ersten Blick der Suche nach der Nadel im Heuhaufen. Die Arzneimittelforschung versucht deshalb, mit wissenschaftlichen Modellen vorherzusagen, welche Moleküle am besten an das jeweilige Ziel-Protein andocken und stark binden. Diese Wirkstoffkandidaten werden anschließend in experimentellen Studien genauer untersucht.

Seit dem Vormarsch der KI nutzt auch die Arzneimittelforschung vermehrt Anwendungen des Maschinellen Lernens. Eine Möglichkeit für solche KI-Anwendungen sind "Graph Neuronale Netze" (GNNs). Sie sollen zum Beispiel vorhersagen, wie stark ein bestimmtes Molekül an ein Zielprotein bindet. Dazu werden GNN-Modelle mit Graphen trainiert, die Komplexe von Proteinen und chemischen Verbindungen (Liganden) darstellen. Graphen bestehen generell aus Knoten, die Objekte darstellen, und Kanten, die Beziehungen zwischen den Objekten anzeigen. In molekularen Graphen von Protein-Liganden-Komplexen gibt es Kanten, die entweder Protein- oder Liganden-Knoten verbinden (und die Struktur von Protein und Ligand erfassen), und andere Kanten, die Protein- und Liganden-Knoten verbinden und spezifische Wechselwirkungen darstellen.

"Wie GNNs zu ihren Prognosen kommen, gleicht einer Black Box, in die man nicht hineinschauen kann", sagt Prof. Dr. Jürgen Bajorath. Der Chemieinformatiker vom LIMES-Institut der Universität Bonn, vom Bonn-Aachen International Center for Information Technology (B-IT) und vom Lamarr-Institut für Maschinelles Lernen und Künstliche Intelligenz in Bonn hat zusammen mit Kollegen der Sapienza Universität in Rom im Detail analysiert, ob die Graph Neuronale Netze tatsächlich Protein-Liganden-Wechselwirkungen lernen, um vorherzusagen, wie stark ein Wirkstoff an ein Zielprotein bindet.

Wie funktionieren die KI-Anwendungen?

Die Forschenden analysierten die Arbeitsweise von insgesamt sechs verschiedenen GNN-Modellen mit ihrer eigens dafür entwickelten "EdgeSHAPer"-Methode und zum Vergleich einer konzeptionell

unterschiedlichen Methode. Diese Computerprogramme “durchleuchten”, ob die GNNs die wichtigsten Interaktionen zwischen Wirkstoff und Protein lernen und damit die Wirksamkeit vorhersagen, wie von Forschern beabsichtigt und erwartet wird – oder ob die KI nur Teil-Prozesse unter die Lupe nimmt und auf anderen Wegen zu den Vorhersagen kommt. “Die GNNs sind sehr abhängig von den Daten, mit denen sie trainiert werden”, sagt der Erstautor der Studie, Doktorand Andrea Mastropietro von der Sapienza Universität in Rom, der für einen Teil seiner Doktorarbeit in der Arbeitsgruppe von Prof. Bajorath in Bonn forschte.

Die Wissenschaftler trainierten die sechs GNNs mit Graphen aus Strukturdaten von Komplexen, für die aus Experimenten bereits die Wirkweise und Bindungsstärke chemischer Verbindungen an ihren Zielproteinen bekannt war. Die trainierten GNNs wurden dann mit anderen Komplexen getestet. Dadurch konnten die Forschenden nachvollziehen, wie die KI funktioniert, um in diesen Berechnungen auf den ersten Blick vielversprechende Vorhersagen zu generieren.

“Wenn die GNNs das machen, was man von ihnen erwartet, müssten sie vor allem die Wechselwirkungen zwischen Wirkstoff und Zielprotein lernen und die Vorhersagen müssten durch Priorisierung spezifischer Wechselwirkungen bestimmt werden”, erläutert Prof. Bajorath. Nach den Analysen des Forscherteams schießen die sechs KI-Programme aber an diesem Ziel vorbei. Die meisten GNNs lernen nur wenige Protein-Wirkstoff-Wechselwirkungen und fokussieren hauptsächlich auf bestimmte Bereiche der Wirkstoffmoleküle. Bajorath: “Um die Bindungsstärke eines Moleküls an ein Zielprotein vorherzusagen, ‘erinnern’ sich die Modelle hauptsächlich an chemisch ähnliche Moleküle, die sie im Training ‘kennengelernt’ haben und an deren Bindungsdaten, unabhängig vom Zielprotein. Diese gelernten chemischen Ähnlichkeiten bestimmen dann im Wesentlichen die Vorhersagen.”

Nach Einschätzung der Wissenschaftler verhält es sich hier weitgehend wie mit dem “Klugen-Hans-Effekt”. Dabei handelt es sich um ein Pferd, das angeblich rechnen konnte. Wie oft Hans mit dem Huf klopfte, sollte das Rechenergebnis mitteilen. Wie sich später erwies, war das Pferd des Rechnens gar nicht mächtig, sondern konnte anhand von Nuancen in Mimik und Gestik seines Begleiters erschließen, um welches Ergebnis es sich handelt.

Was bedeuten die Ergebnisse zur Anwendung dieser Graph Neuronalen Netze für Arzneimittelstudien? “Es ist generell nicht haltbar, dass die GNNs das chemische Zusammenspiel von Wirkstoffen und Proteinen lernen”, stellt der Chemieinformatiker fest. Ihre Vorhersagen sind damit weitgehend überbewertet, weil Prognosen in ähnlicher Qualität mit chemischem Wissen und einfachen Methoden erstellt werden können. Allerdings gibt es auch hier weitere Ansätze für die KI. Zwei der untersuchten GNN-Modelle zeigten eine deutliche Tendenz, mehr Wechselwirkungen zu lernen, wenn die Wirksamkeit bekannter Wirkstoffe zunahm. “Hier lohnt es sich, noch genauer hinzusehen”, sagt Bajorath. Vielleicht ließen sich diese GNNs durch modifizierte Trainingsmethoden weiter in die gewünschte Richtung verbessern. Allerdings müsse man bei der Annahme, dass physikalische Größen auf der Basis molekularer Graphen gelernt werden können, generell vorsichtig sein. „KI ist keine schwarze Magie“, sagt Bajorath.

Noch mehr Licht ins Dunkle der KI

Vielmehr sieht der Chemieinformatiker mit dem open access publizierten EdgeSHAPer und anderen speziell entwickelten Analysetools vielversprechende Ansätze, Licht in die “Black Box” der Künstlichen Intelligenz zu bringen. Der Ansatz seines Teams fokussiert derzeit auf GNNs und neue ‘chemische Sprachmodelle’. “Die Entwicklung von Methoden zur Erklärung von Vorhersagen komplexer Modelle ist ein wichtiges Teilgebiet der KI. Es gibt auch für andere Netzwerkarchitekturen wie Textverarbeitungs-KI Ansätze, die dabei helfen besser zu verstehen, wie das Maschinelle Lernen zu seinen Ergebnissen kommt”, sagt Bajorath. Er erwartet, dass am Lamarr-

Institut, wo er die Rolle eines Chairs für KI in den Lebenswissenschaften übernommen hat, auch im Bereich „Erklärbare KI“ in den nächsten Jahren Spannendes passieren wird.

Publikation: Andrea Mastropietro, Giuseppe Pasculli, and Jürgen Bajorath: Learning characteristics of graph neural networks predicting protein-ligand affinities, Nature Machine Intelligence, DOI: 10.1038/s42256-023-00756-9, Internet: <https://www.nature.com/articles/s42256-023-00756-9>