

Künstliche Intelligenz: Infrastrukturen und Voraussetzungen aus europäischer Perspektive

Artificial Intelligence: Infrastructures and Preconditions from a European Perspective

Sebastian Fritsch, Oliver Maaßen, Morris Riedel

Einleitung

Die Anwendung von Künstlicher Intelligenz (KI) im Rahmen des Maschinellen Lernens ist oft mit der Nutzung großer Datenmengen verbunden, welche leistungsfähige Rechen- und Dateninfrastrukturen benötigen. Gesundheitsdaten stellen hierbei eine besondere Herausforderung dar. Wie es trotzdem gelingen kann, auf europäischer Ebene einen sicheren und gleichzeitig effektiven Umgang mit Daten für die Nutzung in KI aufzubauen, aber auch, wo noch Probleme bestehen, stellt dieser Artikel dar.

Big Data in der Medizin

Box Fallbeispiel

Sie arbeiten seit einiger Zeit auf der Intensivstation Ihres Krankenhauses und haben einen auffälligen Zusammenhang zwischen einer bestimmten Behandlung und einem sehr günstigen Outcome für Ihre Patienten bemerkt. Nach der Auswertung der Daten Ihrer Intensivstation scheint sich Ihre Beobachtung bestätigen. Da Ihnen allerdings die Zahl der behandelten Patienten für eine zuverlässige Bewertung und Einschätzung zu klein vorkommt, suchen Sie nach Möglichkeiten, an weitere Daten zu kommen, in denen Sie nach einem ähnlichen Effekt suchen können. Bei der Suche im Internet stoßen sie auf die Datenbanken "MIMIC", "eICU" und „HiRID“, die Sie unter bestimmten Voraussetzungen nutzen können.

Nutzung von großen Datensätzen

Die Verfügbarkeit ausreichend großer Datensätze in inhaltlich und strukturell guter Qualität ist die zentrale Herausforderung für die Entwicklung der Künstlichen Intelligenz (KI). Die Anwendung von KI im Bereich der Medizin bildet hier keine Ausnahme. Viele nicht-medizinische Forscher haben keinen Zugang zu den besonders geschützten Gesundheitsdaten in Krankenhäusern und Arztpraxen. Doch selbst wenn die Daten verfügbar sind, ist die Anzahl an Datensätzen häufig zu klein oder ihre Qualität (fehlerhafte Messungen, ungenaue/unvollständige Dokumentation, unterschiedliche Datenstrukturen u.ä.) unzureichend. Eine Lösung für dieses Problem bietet die Erstellung von Datenbanken mit anonymisierten, idealerweise multizentrischen Patientendaten aus der klinischen

Routineversorgung. Die Nutzung dieser Daten für nachfolgende Forschungsprojekte wird auch als "secondary usage" bezeichnet.

Box Info:

Die FAIR Data-Prinzipien

Die FAIR Data Prinzipien wurden 2016 von einem Konsortium aus Wissenschaftlern und Organisationen festgelegt und beschreiben wichtige Eigenschaften, die Daten nach ihrer Speicherung erfüllen sollen:

F - Findability
A - Accessibility
I - Interoperability
R - Reusability

Die Daten sollen also wieder auffindbar, langfristig zugänglich, technisch "nachnutzbar" und mit anderen Datensätzen kombinierbar und analytisch und intellektuell wiederverwendbar sein. Zu diesem Zweck sind neben den Informationen, die die Daten selbst enthalten, auch Metadaten unverzichtbar, die Merkmale der Daten beschreiben. Ziel dieser Prinzipien ist eine Verbesserung der maschinellen Auffindbarkeit (d.h. die Fähigkeit von Computersystemen, Daten ohne oder mit nur minimalem menschlichen Zutun zu finden, darauf zuzugreifen und sie wiederzuverwenden), da durch die permanent wachsende Menge und Komplexität von Daten menschliche Nutzer zunehmend auf Unterstützung durch Computer angewiesen sind. [1, 2]

Auch in der medizinischen Forschung wächst das Verständnis für die Notwendigkeit von öffentlich zugänglichen Forschungsdaten (Stichwort: open access), um eine fortgeführte Nutzung zu ermöglichen und so den wissenschaftlichen Fortschritt zu befördern. Besonders in der COVID-19 Pandemie wurde sichtbar, welche Vorteile die Verfügbarmachung von medizinischen Daten haben kann. Gerade in der Frühphase der Pandemie litt die Aussagekraft vieler wissenschaftlicher Veröffentlichungen an der sehr kleinen eingeschlossenen Patientenzahl. Der Zugriff auf zusätzliche, frei verfügbare Daten hätte hier bereits deutlich früher belastbare Ergebnisse liefern können.

Öffentlich verfügbare Datenbanken in der Intensivmedizin

Im Bereich der Intensivmedizin gibt es bereits prominente Beispiele für große Datenbanken, die anonymisierte Patientendaten enthalten: die MIMIC-Datenbanken ("Medical Information Mart for Intensive Care"), die eICU-Datenbank und die HiRID-Datenbank ("high time-resolution ICU dataset").

Die MIMIC-Datenbanken wurden gemeinsam mit dem Massachusetts Institute of Technology am Beth Israel Deaconess Medical Center (BIDMC) in Boston erstellt. Die dritte Auflage, MIMIC-III, enthält intensivmedizinische Daten von über 40.000 Patienten [3]. Erweitert wurde dieses Projekt durch die Veröffentlichung von MIMIC-CXR im Jahr 2019. Diese Datenbank enthält über 375.000 radiologische Aufnahmen inklusive semi-strukturierter Befunde von über 65.000 Patienten aus der Notaufnahme des BIDMC [4]. Im August 2020 wurde schließlich mit MIMIC-IV eine Weiterentwicklung der MIMIC-III Datenbank publiziert [5], die die Anzahl verfügbarer Patienten nochmals auf über 53.000 aufgestockte.

Während die MIMIC-Datenbanken lediglich monozentrisch angelegt sind, umfasst die deutlich umfangreichere eICU Collaborative Research Database [6] multizentrische Daten von insgesamt 139.000 Patienten, die auf über 335 Intensivstationen behandelt wurden. Während die MIMIC- und die eICU-Datenbanken lediglich aus US-amerikanischen Daten

bestehen, enthält die HiRID-Datenbank [7] Daten zu fast 36.000 Patienten der Universitätsklinik für Intensivmedizin des Inselspitals in Bern. Die HiRID-Datenbank zeichnet sich vor allem dadurch aus, dass die Rohdaten der kontinuierlich erhobenen Parameter mit Zeitabständen von 2 Minuten in deutlich höherer Auflösung zur Verfügung stehen als in den anderen Datenbanken.

Man kann erkennen, wie groß der Bedarf an entsprechenden Datenbanken für die Forschung ist, wenn man bedenkt, dass die Erstpublikationen dieser Datenbanken gemeinsam seit Mai 2016, dem Zeitpunkt der Publikation der MIMIC-III Datenbank, bereits mehr als 1200 Mal zitiert wurden.

Die deutsche Medizininformatik-Initiative und die Nationale Forschungsdateninfrastruktur

Das Bewusstsein für die Bedeutung solcher Daten für die klinische Forschung ist auf deutscher und europäischer Ebene in den letzten Jahren deutlich gewachsen. Erkennbar wird dies u. a. an der Medizininformatik-Initiative (MI-I), ein großes, durch das Bundesministerium für Bildung und Forschung mit 180 Mio. Euro gefördertes Programm, welches die Vernetzung von klinischer Versorgung und Forschung verbessern soll (www.medizininformatik-initiative.de/de) [8]. Innerhalb der MI-I haben sich alle deutschen Universitätskliniken mit Partnern aus Forschung und Industrie zu insgesamt vier Konsortien zusammengeschlossen. Wesentliche Elemente der Initiative sind Aufbau und Vernetzung sog. Datenintegrationszentren (DIZ) an den Universitätskliniken. Ziel eines DIZ ist es, die medizinischen Informationen aus den multiplen Systemen der Klinik in standardisierter Form zu sammeln, zu speichern und unter Sicherstellung der rechtlichen Rahmenbedingungen zur Nutzung für die medizinische Forschung bereitzustellen. Das umfasst auch die Erstellung und Wartung von Metadaten, so dass auch „Metadaten-Crawler“, also Programme, die automatisiert nach Metadaten suchen, diese Daten leichter auffinden können. Auf diese Weise werden die technischen und organisatorischen Voraussetzungen für die standortübergreifende Datennutzung zwischen Krankenversorgung und medizinischer Forschung geschaffen, um IT-Lösungen für konkrete medizinische Anwendungen zu entwickeln.

Einen ähnlichen Ansatz verfolgt die aktuell im Aufbau befindliche Nationale Forschungsdateninfrastruktur (NFDI) mit ihrem auf Gesundheitsdaten spezialisierten Teilbereich NFDI4Health (<https://www.nfdi4health.de/>). Letzterer hat das Ziel, eine Forschungsdateninfrastruktur für personenbezogene Gesundheitsdaten aufzubauen, um zur Verschmelzung von epidemiologischer, Public Health- und klinischer Forschung beizutragen.

Box Fallbeispiel

Bei Ihrer täglichen Arbeit auf der Intensivstation nutzen Sie auch die ASIC-App (Algorithmic surveillance of ICU patients with acute respiratory distress syndrome), die im Rahmen der MI-I entwickelt wurde. Während Ihres Dienstes informiert Sie die App auf Ihrem Dienst-Smartphone, welches Sie bei sich tragen, dass einer Ihrer beatmeten Patienten seit der letzten Blutgasanalyse einen Horowitz-Quotienten (p_{aO_2}/FiO_2) von 236 mmHg aufweist. Die App fordert Sie nun auf, auch die weiteren Diagnosekriterien eines ARDS zu bewerten. Als Sie daraufhin das Thorax-Röntgenbild des Patienten betrachten, fällt Ihnen auf, dass dieses nun erstmals beidseitige Infiltrate aufweist. Für eine Herzinsuffizienz oder eine Überwässerung finden Sie hingegen keine Anzeichen. Nachdem Sie diese Angaben in die App eingegeben haben, teilt Ihnen die App mit, dass der Patient die Kriterien eines milden ARDS erfüllt. Sie stellen daher die Diagnose, passen die Beatmungseinstellungen gemäß der entsprechenden Leitlinie an und suchen nach der Ursache.

Datenschutz: Das zweischneidige Schwert

Wenn man über datengetriebene Forschung in der Medizin spricht, kommt man an einem Thema keinesfalls vorbei: dem Datenschutz. Über das allgemeine Für und Wider des Datenschutzes in der Medizin ist viel gesagt und geschrieben worden. Eine umfassende Darstellung würde den Umfang dieses Artikels sicher übersteigen. Trotzdem soll das Thema im Kontext der Nutzung großer medizinischer Datenmengen für die Forschung thematisiert werden.

Führt man sich die aktuelle Situation des Datenschutzes in der EU vor Augen, so kann es nicht verwundern, dass die weiter oben im Text erwähnten Datenbanken nicht aus der EU stammen. Nach Art. 9 Abs. 1 Datenschutz-Grundverordnung (DS-GVO) handelt es sich bei medizinischen Daten um „besondere Kategorien von personenbezogenen Daten“, die der höchsten Schutzstufe unterliegen. Infolgedessen ist ihre Verarbeitung grundsätzlich untersagt, es sei denn, das Verbot wird durch einen besonderen Erlaubnistatbestand aufgehoben. Gesundheitsversorgung und Forschung können theoretisch nach DS-GVO unter bestimmten Umständen eine Verarbeitung ermöglichen. Häufiger ist aber wohl die Einwilligung der oder des Betroffenen. An diese Einwilligung werden jedoch sehr hohe Ansprüche gestellt, die fast schon an eine Einwilligung in einen medizinischen Eingriff erinnern. So werden absolute Freiwilligkeit und eine zuvor erfolgte Aufklärung vorausgesetzt.

Eigentlich ist es nicht vorgesehen, dass eine erlassene EU-Verordnung, die ohne weiteres gesetzgeberisches Verfahren in den Mitgliedsstaaten unmittelbar wirksam wird, nochmals durch die Mitgliedstaaten verändert wird. Genau ist dies im Fall der DS-GVO im Hinblick auf Gesundheitsdaten aber möglich und vor allem in Deutschland auch umfassend geschehen. So gibt es in Deutschland neben der DS-GVO sowohl auf Bundesebene das Bundesdatenschutzgesetz (BDSG) als auch in 13 der 16 Bundesländer weitere spezialisierte Gesetze mit so unterschiedlichen Namen wie Gesundheitsdatenschutzgesetz (NRW), Krankenhausdatenschutzgesetz (Bremen) oder Krankenhausentwicklungsgesetz (Brandenburg), die alle in unterschiedlichem Umfang Datenverarbeitung im Gesundheitswesen regeln. Enthalten die Krankenhausgesetze der Länder selbst keine Vorschriften zum Datenschutz, verweisen sie in der Regel auf die Landesdatenschutzgesetze, welche ggf. auch wiederum auf das BDSG verweisen können. Dies führt zu teilweise skurrilen Situationen: So muss z. B. ein Konsortium der MI-I, also eines deutschlandweit durchgeführten Projektes, für jedes angeschlossene Universitätsklinikum einzeln geprüft werden, welche Gesetze für eine Datennutzung gerade zur Anwendung kommen, welche besonderen Vorschriften gelten und letztlich, welche Form der Datenverarbeitung überhaupt datenschutzkonform an den Standorten möglich ist. Solche Fragen sind ohne Mitwirkung spezialisierter Juristen kaum rechtssicher zu klären.

Die COVID-19-Pandemie hat die Kontroversen beim Thema Datenschutz in Deutschland nochmals deutlich gemacht. So beklagte der Bundesbeauftragte für Datenschutz und Informationssicherheit (BfDI), Prof. Ulrich Kelber, dass in der Pandemie die Grundrechtseingriffe im Bereich des Datenschutzes nicht ausreichend abgewogen und gerechtfertigt seien und dass es hierfür an transparenten Begründungen fehle. Einen deutlich anderen Blick auf die Dinge hat der Bundesverband Informationswirtschaft, Telekommunikation und neue Medien (Bitkom). Dieser zeigte sich irritiert, dass viele Möglichkeiten ungenutzt blieben. Hier spricht man von "datenschutzrechtlicher Prinzipienreiterei", die jene Menschenleben gefährde, die sich durch den flächendeckenden Einsatz digitaler Lösungen retten ließen. Gemeint ist hier vor allem eine Corona-Warn-App mit erweiterter Funktionalität, die allerdings aufgrund der Interventionen von Datenschützern nicht eingesetzt wurden [9].

Besagter Bundesdatenschutzbeauftragte war es auch, der die derzeitige Form der elektronischen Patientenakte für nicht mit der DS-GVO vereinbar hält und den gesetzlichen Krankenkassen mit Maßnahmen droht, falls sie das beschlossene Gesetz wie vorgesehen umsetzen. Genau dies wird aber von den Krankenkassen unter Androhung gesetzlich festgelegter Strafzahlungen verlangt [10]. Diese Beispiele zeigen wie mit einem Brennglas,

wie schwierig das Zusammenspiel von Datenschutz und Gesundheitsversorgung in Deutschland ist. Es bleibt die Frage, ob bei der Abwägung von Datenschutz und Digitalisierung im Gesundheitswesen die Verhältnismäßigkeit immer gewahrt bleibt.

Wenn die Situation in Deutschland schon so kompliziert ist, verwundert es wenig, dass die Umsetzung von Forschungsprojekten innerhalb der EU mit teilweise noch größeren Schwierigkeiten behaftet ist. So kann es passieren, dass trotz der eigentlich EU-weit gültigen DS-GVO bestimmte Arten der Datenverarbeitung in einem Land möglich sind, die in einem anderen hingegen untersagt sind. Leider stellt die Einbindung deutscher Partner in EU-Projekte in Bezug auf den Datenschutz verglichen mit anderen EU-Staaten häufig die größte Herausforderung dar. Für die Planung internationaler Projekte bedeutet dies einen erheblichen Mehraufwand, da schon in einer sehr frühen Planungsphase geprüft werden muss, ob ein geplantes Vorgehen in allen Teilnehmerländern überhaupt möglich ist.

Leider ist nicht zu erwarten, dass sich dieser Zustand der "Kleinstaaterei" in absehbarer Zeit auf EU-Ebene ändern wird. Daher ist es ratsam, Wege zu suchen, um diese Problematik zu umgehen. Ein Beispiel hierfür ist das Föderale Maschinelle Lernen (engl. Federated Machine Learning) oder neuere spezialisierte Verfahren wie „Swarm Learning“, was vor allem auch technische Techniken wie Blockchain und Edge Computing umfasst [11].

Box Hintergrundwissen:

Federated Machine Learning

Möchte man klinikübergreifend mit Daten arbeiten, führt dies leider häufig zu Problemen, viele Kliniken noch immer "Datensilos" darstellen und der Datenschutz eine Herausgabe der Daten nahezu unmöglich macht. Das Prinzip des Federated Machine Learning [12] umgeht dieses Problem, indem verschiedene KI-Modelle vor Ort auf Basis der in den unterschiedlichen Kliniken vorliegenden Daten trainiert und diese Modelle dann in einem zweiten Schritt kombiniert. Diese Modelle enthalten dann keine identifizierbaren Daten, sondern nur mathematische Modelle und deren gelernte Parameter als Ergebnis. Daher greift auch der Datenschutz nicht mehr. Die in den Kliniken trainierten Modelle können dann problemlos zentral zusammengeführt und zu einem finalen Modell berechnet werden. Voraussetzung für den Einsatz dieser Methode ist aber, dass die Kliniken eine ausreichend gut ausgestattete Computing-Infrastruktur vorhalten, mit deren Hilfe KI-Modelle erzeugt und lokal berechnet werden können.

Von Beginn an war man sich auch in den Reihen der MI-I der zu erwartenden Schwierigkeiten mit dem Datenschutz bewusst. Darum war es auch ein zentrales Ziel der Initiative, Grundlagen dafür zu schaffen, dass Austausch und Nutzung der Daten aus Krankenversorgung, klinischer und biomedizinischer Forschung über die Grenzen von Institutionen und Standorten hinweg in datenschutzkonformer Weise ermöglicht wird. Da die in den einschlägigen Gesetzen benannten Erlaubnistatbestände für eine Umsetzung der MI-I in vollem Umfang nicht ausreichen, ist die informierte Einwilligung des Patienten der einzige Ausweg, was eine Reihe neuer Probleme aufwirft.

Laut DS-GVO ist jede Datenverarbeitung an einen konkreten Zweck gebunden. Das bedeutet im Forschungskontext aber auch, dass ein Patient der Nutzung seiner Daten nur zustimmen kann, wenn ihm das Ziel des Forschungsprojektes bekannt ist. Ebenso dürfen die Daten nur so lange gespeichert werden, wie der Studienzweck es erfordert. Anschließend müssen die Daten entweder gelöscht oder so vollständig anonymisiert werden, dass eine Re-Identifikation der Patienten ausgeschlossen ist. Da hierbei aber immer auch relevante Informationen verloren gehen, schränkt dieses Vorgehen die weitere Nutzbarkeit medizinischer Daten erheblich ein.

Es liegt in der Natur der Sache, dass große Datenbanken medizinischer Daten für unterschiedlichste Ziele verwendet werden, wie ganz praktisch ein Blick in die Zitierung der MIMIC-Daten erkennen lässt. Daher ist es schlichtweg nicht möglich, bereits jetzt alle erdenklichen Anwendungsbereiche für die im Rahmen der MI-I erhobenen Daten zu beschreiben. Dies macht es erforderlich beim Patienten eine breite Zustimmung (sog. "Broad Consent") einzuholen, die die Nutzung der Daten möglich macht. Die ersten Erfahrungen des Arbeitskreises Consent der MI-I zeigen bereits, dass es sich hierbei um ein hochkomplexes Thema mit ganz unterschiedlichen Schwerpunkten handelt, so dass eine erfolgreiche Umsetzung des Broad Consent keineswegs sicher ist. So wurde bereits jetzt klar, dass die verantwortlichen Stellen große Vorbehalte gegen eine Kopplung der Aufklärung über die Datennutzung durch die MI-I an die administrative Aufnahme im Krankenhaus haben. Eine Trennung dieser beiden Aufklärungen würde aber zu einem erheblichen zusätzlichen Personalaufwand und unweigerlich auch zu Verlusten von Patientendaten führen, wenn bei diesen eine Einwilligung aus unterschiedlichsten Gründen nicht mehr eingeholt werden kann.

Lösungen für die beschriebene Problematik sind zwar denkbar, aber wohl kaum ohne breite gesellschaftliche und politische Diskussionen möglich. Denkbar wäre einerseits das Konzept der Datenspende, bei der Personen bereits vor dem Auftreten einer Erkrankung in die Nutzung ihrer Gesundheitsdaten einwilligen können. Vorteilhaft daran wäre, dass die Entscheidung mit ausreichend Bedenkzeit und unabhängig von der akuten Erkrankungssituation treffen könnten. Denkbar wäre auch eine Etablierung deutscher Online-Portale, wie sie in den USA bereits existieren (<https://www.patientslikeme.com/>), in denen betroffene Patienten ihre Erfahrungen, Symptome und Daten austauschen und die Möglichkeit haben, ihre Daten unkompliziert für eine Nutzung im Forschungskontext freizugeben.

Eine theoretische Alternative stellt die Änderung der gesetzlichen Grundlagen dar, so dass unter Wahrung höchster Datenschutzstandards (und entsprechender Strafbewehrung bei Verstößen) eine Datennutzung für Wissenschaft und Qualitätssicherung auch ohne explizite Einwilligung möglich gemacht würde.

Es soll an dieser Stelle nicht der Eindruck vermittelt werden, dass der Datenschutz im Gesundheitswesen in erster Linie als „Klotz am Bein“ zu sehen ist. Es ist absolut unbestritten, dass solch hochsensible Daten wie Gesundheitsdaten zu keiner Zeit in falsche Hände geraten dürfen. Auch sind uns Vorbehalte in der Bevölkerung und der Ärzteschaft gegen eine allzu liberale Auslegung des Datenschutzes durchaus bewusst. Dennoch plädieren wir dafür, dass bei der Debatte stärker die Verhältnismäßigkeiten in den Blick genommen werden. Insbesondere sollte auch offen darüber diskutiert werden, welche negativen Folgen eine rigide Auslegung des Datenschutzes für die Bürger haben kann. Im Rahmen der Diskussionen um die Corona-App hat Alena Buyx, die Vorsitzende des Deutschen Ethikrates, diese Problematik auf den Punkt gebracht: "Stärkeren Datenschutz in der Umsetzung, das gibt's kaum. [...] Und da muss man so langsam wirklich fragen - und da bin ich nicht die Einzige - wir schränken so viele Grundrechte ein und den Datenschutz so gar nicht? Bei dem sind wir bis auf Punkt und Komma extrem präzise?" [13]. Diese Frage lässt sich analog sicherlich auch auf die Situation der Forschung mit Gesundheitsdaten übertragen.

Datenanalyse mit Methoden der KI

Box Fallbeispiel

Zu Ihrer Freude stellen Sie fest, dass die Datenbanken, die sie sich näher ansehen, Daten enthalten, die sie auf Ihre Fragestellung hin auswerten können. Ein befreundeter Informatiker, dem Sie von Ihrem Projekt erzählt haben, gibt Ihnen den Tipp, dass Sie vielleicht noch bessere Ergebnisse erhalten, wenn Sie die Daten nicht ausschließlich mit

statistischen Methoden, sondern auch mit KI-basierten Modellen auswerten. Er bietet Ihnen an, Ihnen bei der rechenintensiven Auswertung mit besagten KI-Modellen zu helfen. Da er Zugriff auf einen Supercomputer hat, beschließen Sie, auch diese Ressource zu nutzen, um die Auswertung erheblich zu beschleunigen.

Erstellung von KI-Modellen

Auch wenn die Datenanalyse mit KI-Modellen teilweise statistische Verfahren verwendet, unterscheidet sie sich doch von rein statistischen Ansätzen. Grob vereinfachend kann man sagen, dass Statistik in erster Linie versucht, vorliegende Daten mithilfe von Wahrscheinlichkeitstheorien möglichst gut zu analysieren und zu beschreiben. Die Nutzung von KI zielt hingegen darauf ab, aus gegebenen Daten zu lernen und Vorhersagen über Zukünftiges zu treffen. Auch im prinzipiellen Vorgehen unterscheiden sich KI und Statistik und haben so ihre je eigenen Stärken. So ist in der Statistik die Formulierung der Fragestellung und die Auswahl eines geeigneten Studiendesigns von besonderer Bedeutung. Auch können Aspekte der Fallzahlplanung und der Kontrolle möglicher Datenverzerrungen (Bias) berücksichtigt werden. Bei der Nutzung der KI steht gemeinhin der Bereich der Datenanalyse deutlich stärker im Vordergrund, da KI stärker explorativ arbeitet. Dies zeigt sich beispielsweise auch an der schon erwähnten Nutzung von Sekundärdaten, die eigentlich für andere Zwecke als die durchgeführte Analyse erhoben wurden. Abweichend von der Statistik, in der Analysen meist nur singulär durchgeführt werden, nutzt KI im Rahmen des maschinellen Lernens (engl. Machine learning) iterative Lernverfahren, die auf komplexen Optimierungsalgorithmen beruhen, wie z.B. Stochastic Gradient Descent (SGD) oder Adam. Das bedeutet, dass ein einzelner Datensatz durchaus mehrfach analysiert wird, um ein Muster in den Daten ("pattern") zu identifizieren (z.B. mehrere Epochen und Batches beim Deep learning). Letztendlich lassen sich KI und Statistik dennoch nur unvollständig trennen, da weite Teile von KI auf statistischer Lerntheorie beruhen (i.e.S. Probably Approximately Correct Learning, „PAC Statement“).

Um KI überhaupt sinnvoll anwenden und erfolgreiche KI-Modelle entwickeln zu können, sind Datensätze von guter Qualität von zentraler Bedeutung. Die Datensätze sollten möglichst vollständig, korrekt bzw. konsistent und unverzerrt sein. Ist die Qualität der Inputdaten für das KI-Modell nicht gut, so führt dies häufig auch zu einer schlechten Vorhersagefähigkeit des Modells (engl. "rubbish-in, rubbish-out"). Auch wenn keine pattern in den Daten erwartet werden, ist die Anwendung eines ML-Modell zumeist nicht sinnvoll.

Das KI-Modell selbst wird durch verschiedenartige Datenanalyse-Prozesse erzeugt. Die Inputdaten-Mengen werden zuerst vorbereitet, indem man sich auf gewisse Datenmerkmale konzentriert (engl. "Feature Selection/Engineering"), die im konkreten medizinischen Anwendungsfall relevant sind. In diesem Prozess ist es auch hilfreich schon ein KI-Modell aus der Vielzahl möglicher Modelle auszusuchen (z.B. Entscheidungsbäume, Neuronale Netze, Support Vector Machines, etc.), da in bestimmten Fällen die Daten für unterschiedliche KI-Modelle in spezieller Form kodiert werden müssen. In dem Kontext „Feature Learning“ wurde vor allem durch Anwendung tiefer Neuronaler Netze (engl. deep learning) erhebliche Durchbrüche erzielt da „Features“ automatisch gelernt werden, statt vom Menschen manuell vorbereitet zu werden. Darüber hinaus haben tiefe neuronale Netze auch spezifische Modellarchitekturen, die auf bestimmte Datenstrukturen (z.B. Bilddaten, Zeitreihendaten) sehr effektiv angewendet werden können. Das Lernen tiefer Neuronaler Netze ist daher eine spezifische Unterklasse von Maschinellen Lernmodellen, die in den letzten Jahren enorm an Bedeutung im Zusammenhang mit großen Datenmengen gewonnen hat.

Die Vorbereitung beinhaltet auch das Aufteilen der gesamten Datensätze in einen Trainings- und einen Testdatensatz. Sollten sehr viele Daten verfügbar sein, ist es ratsam, zusätzlich einen oder mehrere Validierungsdatsätze zu erzeugen, um später eine bessere Performance bei der Selektierung des Modells zu erreichen. Hat man das passende Modell

ausgewählt, die Rahmenbedingungen angepasst und seine Daten erfolgreich vorbereitet, ist es sinnvoll, diese Parameter bereits als offene Datensätze unter den FAIR Prinzipien zu teilen. So erreicht man ein einheitliches Vorgehen und erspart anderen Klinikern Zeit bei der Vorbereitung ihrer eigenen Datensätze.

Der Trainingsdatensatz wird nun mit dem Algorithmus eines ausgewählten KI-Modells untersucht, was man als "Training" des KI-Modells bezeichnet. Der Output dieses Prozesses ist ein KI-Modell mit Schätzern, die oft durch einen Optimierungsalgorithmus gefunden werden. Im folgenden Schritt wird die Genauigkeit des trainierten KI-Modells auf der Basis des Testdatensatzes, also der Daten, die nicht zum Training benutzt worden sind, überprüft (engl. "model generalization"). Sollten Validierungsdatensätze zur Verfügung stehen, werden diese zur Parametertuning, d.h. zur Bestimmung sogenannter Hyperparameter, die das Lernverhalten der KI-Algorithmen kontrollieren und nicht durch Optimierungsroutinen gelernt werden oder zur Modellordnung verwendet. Dies stellt eine der größten Herausforderungen im maschinellen Lernen dar. Hat man ein KI-Modell mit guter Genauigkeit erstellen können, kann auch das Modell selbst, zum Beispiel im Rahmen des Federated Machine Learnings, geteilt werden. Möchte man Hochleistungsrechnen (engl. High performance computing, HPC) nutzen, kann es auch sehr sinnvoll sein, sog. HPC-Schedulingskripte und Datenanalyse-Skripte online zu teilen, da Hochleistungsrechner in Deutschland Benutzern oft Zugriff nach gleichen Prinzipien gewähren (engl. "scheduling").

Box Info:

Das "Handle-System"

Es bietet sich an, die Modelle und Datensätze online mit Systemen zu teilen, die nach dem "Handle-System" (engl. für Handgriff) funktionieren. Mithilfe dieses Systems ist es möglich, eine permanente Verbindung zwischen einem eindeutigen und dauerhaften persistenten Identifikator (engl. Permanent Identifier, PID) und den zugehörigen Daten zu erzeugen. Mithilfe des PID kann eine dauerhafte Lokalisierbarkeit des Datensatzes unabhängig vom physischen Speicherort erreicht werden. Der Datensatz wird auf diese Weise referenzierbar, d. h. es ist möglich, die "Handles" von Daten und Modellen in Publikationen oder Anträge einzubinden. Das wohl bekannteste Handle-System in der Medizin ist der Digital object identifier (DOI), der z. B. zur Identifikation wissenschaftlicher Publikationen genutzt wird.

Hochleistungsrechnen (High Performance Computing)

Die meisten hier erwähnten Prozesse nutzen Ansätze des Hochleistungsrechnens (engl. "High Performance Computing", HPC). Als Hochleistungscomputer bezeichnet man dabei grundsätzlich einen für seine Zeit sehr schnellen Computer. Die Bauweise des Computers ist hierbei zunächst völlig unerheblich (Abb. 1).



Abb. 1: Mit JUWELS verfügt das Forschungszentrum Jülich über einen Supercomputer, der zu den schnellsten der Welt gehört. Quelle: Forschungszentrum Jülich.

Die zentrale Domäne des HPC ist die Berechnung, Modellierung und Simulation komplexer Systeme und die Verarbeitung riesiger Messdatenmengen. Während die HPC in Wissenschaftsdisziplinen wie Wetter- und Klimaforschung, Astro- und Teilchenphysik, Quantenchemie und Strömungsmechanik bereits seit langer Zeit fest etabliert ist, gewinnt sie in der Medizin erst langsam an Bedeutung. Gründe für diese zögerliche Entwicklung sind neben der bereits beschriebenen schlechten Verfügbarkeit der notwendigen Daten sicher auch die hohen Kosten, die Aufbau und Betrieb eines solchen Systems mit sich bringt. Es ist daher wahrscheinlich, dass der Einsatz von HPC lediglich an einzelnen Universitätskliniken und dann nur in kleinerem Maßstab erfolgen wird. Beispiele für die HPC Nutzung in der Medizin sind die digitale Bildverarbeitung in Radiologie [14] und Pathologie [15] oder auch die Analyse von kontinuierlich erhobenen Vitalparametern in OP und Intensivstation [16].

Hält man sich vor Augen, dass moderne Hochleistungsrechner, die in der oberen Liga der Leistungsfähigkeit mitspielen, Investitionskosten im Bereich von hohen zweistelligen, oft dreistelligen Euro-Millionenbeträgen mit sich bringen, ist klar, dass man daran interessiert ist, die Nutzung eines solchen Computers effizient zu koordinieren. Während diese Koordinierung in Deutschland durch das Gauss Centre for Supercomputing erfolgt, unterstützt seit 2007 auf europäischer Ebene die Initiative "Partnership for Advanced Computing in Europe" (PRACE) die Bündelung der HPC-Kapazitäten. Diese Bemühungen wurden in den vergangenen Jahren durch die Europäische Union durch die Gründung von "European High Performance Computing" (EuroHPC) nochmals intensiviert. So hat die EU-Kommission im September 2020 beschlossen, die Mittel für EuroHPC auf 8 Milliarden Euro (verteilt auf 13 Jahre) aufzustocken, um dafür zu sorgen, dass das europäische Hochleistungsrechnen einen Platz in der Weltspitze behaupten kann [17]. Auch in Deutschland hat man der zunehmenden Rolle der HPC dadurch Rechnung getragen, dass man die Förderung im Rahmen des Verbundes für Nationales Hochleistungsrechnen (NHR) ausgebaut hat [18]. Ziel soll es sein, den Wissenschaftlern der deutschen Hochschulen die für ihre Forschung benötigte Rechenkapazität zur Verfügung zu stellen und ihre Kompetenzen zur effizienten Nutzung dieser Ressource zu stärken. Insbesondere bei der Nutzung von KI-Algorithmen bieten HPC-Infrastrukturen einen enormen Vorteil, um bspw. das Erlernen tiefer künstlicher Neuronaler Netze erheblich zu beschleunigen.

Box Fallbeispiel

Nachdem sie alle notwendigen Vorbereitungen getroffen haben (u. a. positives Ethikvotum und Zustimmung des Datenschutzbeauftragten, Anonymisierung des Datensatzes), möchten Sie dem Informatiker, der nun mit Ihnen zusammenarbeitet, gerne die auszuwertenden Daten zur Verfügung stellen. Jedoch ist Ihnen der beste Weg dazu noch nicht klar. Bei einer entsprechenden Internetrecherche stoßen Sie auf den Dienst "B2DROP" der EUDAT Collaborative Data Integration, mit dem es möglich ist, Ihre Daten problemlos zu speichern und mit anderen Forschern sicher auszutauschen. Sie einigen sich, die Daten im Anschluss an den Transfer bei B2SHARE dauerhaft und referenzierbar zu hinterlegen. Sie benutzen damit Dienste der European Open Science Cloud (EOSC).

Die European Open Science Cloud (EOSC)

Ein zentraler, wenn nicht sogar der Baustein für die datengetriebene Forschung in Europa stellt die European Open Science Cloud (EOSC) dar. Die EOSC ist ein Projekt der EU, mit dessen Hilfe europäischen Wissenschaftlern der Zugang zu wissenschaftlichen Daten, Plattformen und Dienstleistungen für die Datenverarbeitung erleichtert werden soll. Ziel ist es, 72 bestehende Forschungsdateninfrastrukturen in Europa zu vereinigen und ein Netz von FAIR-Daten und zugehörigen Diensten für die Wissenschaft zu realisieren, das Forschungsdaten gemäß den FAIR-Prinzipien interoperabel und maschinenverarbeitbar macht. Für den Bereich der Life Sciences wurde mit EOSC LIFE (<https://www.eosc-life.eu/>) eine spezialisierte Untereinheit gebildet, die die besonderen Herausforderungen aus diesem Bereich (z.B. den Datenschutz) in den Blick nimmt. Flankiert wird das Projekt von den europäischen Infrastrukturen wie EGI, EUDAT, GÉANT und OpenAIRE, die europaweit und darüber hinaus umfassende Datenmanagementdienste zur Nutzung in der Wissenschaft bereitstellen. Unter diesen ist die EUDAT Collaborative Data Infrastructure (EUDAT CDI) eine der größten Infrastrukturen von integrierten Datendiensten und Ressourcen zur Unterstützung der Forschung in Europa. Sie wird von einem Netzwerk von mehr als 20 europäischen Forschungsorganisationen, Daten- und Rechenzentren getragen. Konzeptionelle Aspekte wie Best-practices, Guidelines und Definition zu PID, FAIR und Datasharing im Allgemeinen werden auf der europäischen Ebene von der Research data alliance (RDA, <https://rd-alliance.org>) in unterschiedlichen Gruppen entwickelt. Die verschiedenen Gruppen der RDA decken ein breites wissenschaftliches Spektrum ab, u.a. auch die biomedizinischen Wissenschaften (<https://rd-alliance.org/rda-disciplines/rda-and-biomedical-sciences>).

Die EUDAT CDI bietet mit den sogenannten B2-Services eine Reihe von Dienstleistungen an, die den Umgang mit Forschungsdaten deutlich erleichtern. Einer dieser Dienste ist B2DROP, der es Wissenschaftlern ermöglicht, ihre Forschungsdaten zu speichern, mehrere Datenversionen mit anderen Forschern zu synchronisieren und aktuell zu halten und schließlich auch austauschen zu können. Während bei B2DROP der Austausch mit Kollegen im Vordergrund steht, stellt die Ablage der Daten bei B2SHARE eine echte Veröffentlichung der Daten dar, welche im Rahmen der FAIR Data-Prinzipien z.B. von Journals vor einer Veröffentlichung eines Forschungsartikels zunehmend gefordert wird. B2SHARE bietet dabei als vollwertige Repository bei der Ablage der Daten neben der Vergabe eines dauerhaften Identifikators (Digital object identifier, DOI) auch die Verknüpfung des Datensatz mit Metadaten. Ein Beispiel eines medizinischen Datensatzes in B2SHARE findet sich unter folgendem Link: <https://b2share.fz-juelich.de/records/aef5d3b8aa044485b9620b95b60c47a2>.

Die hinterlegten Metadaten erhöhen die Wahrscheinlichkeit, dass der Datensatz von anderen Forschern gefunden, ein weiteres Mal verwendet und anschließend natürlich auch zitiert wird. Um mithilfe dieser Metadaten weitere interessante Datensätze zu finden, steht mit B2FIND ein weiteres Tool zur Verfügung. Bei B2FIND handelt es sich um einen sogenannten Metadata harvester, der zahlreiche Repositories durchsucht und die

Metadaten "erntet", um sie für Suchvorgänge zugänglich zu machen. Zu den Metadaten zählen zum Beispiel eine Beschreibung der eingeschlossenen Patientenpopulation oder der Zeitraum der Datensammlung. Hat man selbst einen Datensatz mit ähnlicher Population, kann man so nach weiteren entsprechenden Datensätzen suchen und sie gemeinsam in die Analyse einbeziehen. Das Suchergebnis, welches B2FIND erzielt, wenn man nach dem oben aufgeführten Beispieldatensatz sucht, ist unter folgendem Link einzusehen: <http://b2find.eudat.eu/dataset/efe7389a-d5b6-570c-a956-df25eab6d9bb>.

The way from bench to bed

Box Fallbeispiel

Nachdem Sie Ihre Analyse durchgeführt haben, stellen Sie fest, dass das Ergebnis leider nicht so einfach interpretierbar ist, wie sie es zuvor erwartet haben. Zwar lässt sich der positive Effekt, den Sie während Ihrer Arbeit auf der Intensivstation beobachtet haben, tatsächlich nachweisen, jedoch zeigt die Analyse mit KI-basierten Methoden, dass die Daten mehrere Cluster mit unterschiedlichen Ergebnissen bilden. Dies deuten sie so, dass zusätzliche Parameter mit in die Analyse einbezogen werden müssen, damit Ihre Patienten zuverlässig von dem gefundenen Effekt profitieren können. Dies würde in einem recht komplizierten Score resultieren, so dass Sie befürchten, dass ihn in der täglichen Routine niemand anwenden würde. Sie überlegen daher, ob sie ein Software-Programm schreiben, welches den Score selbständig berechnet und Empfehlungen zur Therapie Ihrer Patienten gibt.

In den letzten Jahren wurden mehr und mehr mögliche Anwendungen von KI in Anästhesie und Intensivmedizin publiziert, aber bisher hat es kaum ein Produkt es in die routinemäßige Anwendung am Patientenbett geschafft. Die Nutzung von Software im Allgemeinen und damit auch von KI-Modellen bei der Behandlung von Patienten unterliegt dem Medizinprodukterecht. Mit einer pandemiebedingten Verzögerung von einem Jahr ist am 26. Mai 2021 die EU-Medizinprodukte-Verordnung (EU) 2017/745 (Medical Device Regulation, MDR) [19] in Kraft getreten und hat damit das deutsche Medizinproduktegesetz (MPG) weitgehend abgelöst. Ziel der MDR ist primär der Schutz des Patienten vor fehlerhaften oder risikobehafteten Medizinprodukten. Mit der neuen Verordnung geht eine Ausweitung des Gültigkeitsbereiches und eine deutlich strengere Prüfung vor der Markteinführung einher. Ebenso neu ist die Verpflichtung der Hersteller zur Überwachung ihrer Produkte nach der Markteinführung (Post market surveillance). Die MDR teilt Medizinprodukte allgemein in vier Risikoklassen ein (I, IIa, IIb, III), die den Umfang der Überprüfung und Überwachung bestimmen.

Box Rechtliches:

Risikoklassen von Software als Medizinprodukt

Medical Device Regulation, Anhang VIII, Kapitel III, 6.3. Regel 11:

Software, die dazu bestimmt ist, Informationen zu liefern, die zu Entscheidungen für diagnostische oder therapeutische Zwecke herangezogen werden, gehört zur Klasse IIa, es sei denn, diese Entscheidungen haben Auswirkungen, die Folgendes verursachen können:

- den Tod oder eine irreversible Verschlechterung des Gesundheitszustands einer Person; in diesem Fall wird sie der Klasse III zugeordnet, oder
- eine schwerwiegende Verschlechterung des Gesundheitszustands einer Person oder einen chirurgischen Eingriff; in diesem Fall wird sie der Klasse IIb zugeordnet.

Software, die für die Kontrolle von physiologischen Prozessen bestimmt ist, gehört zur Klasse IIa, es sei denn, sie ist für die Kontrolle von vitalen physiologischen Parametern

bestimmt, wobei die Art der Änderung dieser Parameter zu einer unmittelbaren Gefahr für den Patienten führen könnte; in diesem Fall wird sie der Klasse IIb zugeordnet. Sämtliche andere Software wird der Klasse I zugeordnet.

Medizinprodukte der Klasse I können mit einem durch den Hersteller selbst durchgeführten Konformitätsverfahren die CE-Kennzeichnung erhalten, die eine Markteinführung erst ermöglicht. Bei allen anderen Risikoklassen ist ein Konformitätsverfahren durch sogenannte Benannte Stellen erforderlich. Grundsätzlich muss ein Medizinprodukt seine Leistung und Sicherheit auf der Basis klinischer Daten nachweisen. Während es aber früher im Rahmen des MPG durchaus möglich war, hierfür ausschließlich Daten aus der Literatur oder aus der Testung ähnlicher Produkte zu verwenden, so ist dies seit Einführung der MDR nur noch in sehr beschränktem Umfang möglich. So müssen die herangezogenen Daten mittlerweile aus einem Test eines nahezu identischen Produkts, das auch nahezu identisch verwendet wurde, stammen (sog. Äquivalenzprinzip). Andernfalls müssen die erforderlichen Daten durch klinische Prüfungen, umgangssprachlich "Studien" genannt, geliefert werden. Bei Medizinprodukten der Klasse III sind diese klinischen Prüfungen inzwischen aufgrund der strengen Regularien praktisch der Regelfall.

Es erklärt sich von selbst, dass die Testung eines bisher nicht zertifizierten Medizinproduktes mit hohen Anforderungen an die Sicherheit der Patienten bzw. Probanden verbunden ist. Ihre Durchführung ist daher entsprechend aufwendig und auch teuer. Gleichzeitig fällt wohl die Mehrheit aller denkbaren Anwendungen von KI in der Anästhesie und Intensivmedizin schon aufgrund des instabilen Zustands der dort behandelten Patienten in eine hohe bis sehr hohe Risikoklasse. Neben der Durchführung der klinischen Prüfung entstehen auch durch die Einbindung der Benannten Stelle zusätzliche Kosten, bevor ein Einsatz am Patienten möglich ist.

Dies alles führt letztendlich zu der bedauerlichen Situation, dass neue Entwicklungen aus der Forschung trotz erfolgversprechender Ansätze nicht wehr den Weg in die Praxis finden. Ändern ließe sich dieser Umstand nur auf gesetzgeberischer Ebene. Allerdings ist aufgrund der Tendenz, die Gesetzgebung im Bereich der Medizinprodukte eher zu verschärfen, damit wohl vorerst nicht zu rechnen.

Eine Möglichkeit, die Sicherheit von Patienten nicht zu gefährden und die Anwendung von KI-basierter Software trotzdem zu evaluieren, stellt entweder die Nutzung simulierter Daten oder aber - und so schließt sich der Kreis - die Nutzung großer Datenbanken, die auch Vitalparameter als sogenannten Wave form-Data enthalten, dar. Die Daten könnten der Software pseudo-prospektiv zur Verfügung gestellt werden, d.h. sie werden "abgespielt" als ob sie aus dem Echtzeit-Monitoring eines Patienten stammen. Die Software hätte damit lediglich die Daten zur Verfügung, die bis zu einem bestimmten Zeitpunkt erhoben wurden und nicht den gesamten Datensatz. Die Reaktionen der Software würden dann beobachtet und ärztlich evaluiert. Diese Art der Überprüfung kann natürlich keine klinische Studie ersetzen, da die in ihr die Nutzer-Sichtweise (Mensch-Software-Interaktion) völlig ausgeklammert ist. Aber zumindest die grundlegende technische Funktionalität könnte so getestet werden, was den Umfang dessen, was dann noch real am Patientenbett zu testen ist, deutlich reduzieren würde. Leider sehen die entsprechenden Verordnungen eine solche Möglichkeit bisher nicht vor.

Zusammenfassung

Auch wenn es eine Vielzahl ermutigender Berichte über die Anwendung von Künstlicher Intelligenz in der Anästhesie und Intensivmedizin gibt, so steht die praktische Anwendung doch noch immer vor großen Herausforderungen. In diesem Artikel haben wir anhand einer Anwendung von KI in der Anästhesie und Intensivmedizin dargestellt, wo Hindernisse für den Fortschritt liegen, aber auch wie zahlreiche Initiativen dazu beitragen, diese

Hindernisse anzugehen und diese aus dem Weg zu räumen. Ein überaus wichtiger Punkt stellt hierbei die Gesetzgebung dar, die in weiten Teilen noch nicht auf die Herausforderungen, die der Einsatz von KI in der Medizin mit sich bringt, reagiert hat. Hier gilt es, in der Zukunft Aufklärungsarbeit zu leisten, damit die Besonderheiten der Künstlichen Intelligenz bei weiteren Gesetzesänderungen stärker berücksichtigt werden können. Erfreulicherweise tritt zunehmend der Umstand ins Bewusstsein, welche Bedeutung die Integration und Vernetzung von Daten zu deren Nutzung für den technischen Fortschritt haben. Wenn dieser eingeschlagene Weg konsequent weiter beschritten wird, wird ein sinnvoller Einsatz von Künstlicher Intelligenz in der Medizin endlich auch in Europa möglich.

Kernaussagen

- Medizinische Forschungsdaten sollen den FAIR Data-Prinzipien entsprechen, also wiederauffindbar, langfristig zugänglich, mit anderen Systemen kompatibel und wiederverwendbar sein.
- Für die intensivmedizinische Forschung gibt es bereits mehrere große Datenbanken mit anonymisierten Patientendaten, die praktisch frei genutzt werden können.
- Der Datenschutz wird bis hinunter auf Länderebene teilweise unterschiedlich gesetzlich geregelt, was die Durchführung multizentrischer Projekte mitunter erschweren kann.
- Eine einheitliche, rechtskonforme Umsetzung einer umfassenden Einwilligung („Broad Consent“) an deren Ende die Möglichkeit zur umfassenden Nutzung von Gesundheitsdaten steht, ist insbesondere im Bereich der KI sehr wünschenswert.
- Für die Entwicklung von KI-Modellen müssen Datensätze von möglichst hoher Qualität sein; ist diese unzureichend beeinflusst dies die Performance eines KI-Modells negativ.
- Für viele Anwendungen der KI sind Hochleistungsrechner erforderlich, so dass derzeit große Summen in diesen Bereich investiert werden.
- Der Aufbau einer umfassenden europäischen Forschungsdateninfrastruktur, der European Open Science Cloud (EOSC), welche zahlreiche bereits bestehende Strukturen bündeln und integrieren soll, stellt ein zentrales Element für die europaweite Verfügbarmachung von Forschungsdaten nach den FAIR Data-Prinzipien dar.
- Software, die im medizinischen Bereich angewendet wird, kann unter die Medical device regulation (MDR) fallen, welche die Vorgaben für die Marktzulassung von Medizinprodukten deutlich verschärft hat.
- Die aufwendigen Verfahren zur Marktzulassung führen dazu, dass viele gute Ansätze aus der Forschung mit KI nicht in der täglichen Routine am Patientenbett ankommen.

Zu den Autoren

Sebastian Fritsch

Dr. med., Facharzt für Anästhesiologie, 2014–2019 Facharztausbildung am Universitätsklinikum der RWTH Aachen. Seit 2014 wissenschaftlicher Mitarbeiter in der Klinik für Operative Intensivmedizin und Intermediate Care am Universitätsklinikum der RWTH Aachen. Seit 2020 wissenschaftlicher Mitarbeiter am Juelich Supercomputing Center des Forschungszentrums Jülich.

Oliver Maaßen

M.Sc. Public Health, 2012 - 2016 Projektleiter Healthcare IT, Nexus AG, seit 2016 wissenschaftlicher Mitarbeiter und Projektleiter in der Klinik für Operative Intensivmedizin und Intermediate Care am Universitätsklinikum. Seit Projektstart des SMITH Projektes am 01.01.2018 Standortkoordinator in der Uniklinik RWTH Aachen für das SMITH Projekt und Koordinator des intensivmedizinischen Use Case ASIC.

Morris Riedel

Prof. Dr.-Ing., Full Professor an der School of Engineering and Natural Sciences der University of Iceland, Reykjavik. Seit 17 Jahren Tätigkeiten im Bereich parallel and distributed systems. Weiterhin Tätigkeit am Juelich Supercomputing Centre des Forschungszentrums Jülich als Leiter mehrerer Forschungsgruppen mit den Schwerpunkten „High Productivity Data Processing“ und „Deep Learning“.

Korrespondenzadresse:

Dr. med. Sebastian Fritsch
Forschungszentrum Jülich GmbH
Jülich Supercomputing Center
52425 Jülich
s.fritsch@fz-juelich.de

Referenzen

1. Wilkinson MD, Dumontier M, Aalbersberg IJ et al. The FAIR Guiding Principles for scientific data management and stewardship. *Scientific Data* 2016; 3: 160018. DOI: 10.1038/sdata.2016.18
2. Directorate-General for Research and Innovation (European Commission). Realising the European open science cloud. First report and recommendations of the Commission high level expert group on the European open science cloud 2016. DOI: 10.2777/940154
3. Johnson AE, Pollard TJ, Shen L et al. MIMIC-III, a freely accessible critical care database. *Scientific data* 2016; 3: 1-9.
4. Johnson AEW, Pollard TJ, Berkowitz SJ et al. MIMIC-CXR, a de-identified publicly available database of chest radiographs with free-text reports. *Sci Data* 2019; 6: 317. DOI: 10.1038/s41597-019-0322-0
5. Johnson A, Bulgarelli L, Pollard T et al. MIMIC-IV (version 1.0). *PhysioNet* 2021. DOI: 10.13026/s6n6-xd98.
6. Pollard TJ, Johnson AE, Raffa JD et al. The eICU Collaborative Research Database, a freely available multi-center database for critical care research. *Scientific data* 2018; 5: 1-13.
7. Hyland SL, Faltys M, Huser M et al. Early prediction of circulatory failure in the intensive care unit using machine learning. *Nat Med* 2020; 26: 364-373. DOI: 10.1038/s41591-020-0789-4
8. Semler SC, Wissing F, Heyder R. German Medical Informatics Initiative. *Methods Inf Med* 2018; 57: e50-e56. DOI: 10.3414/ME18-03-0003
9. MDR. Kritik am Umgang mit Datenschutz in der Corona-Pandemie (25.03.2021) Im Internet: <https://www.mdr.de/nachrichten/deutschland/politik/corona-datenschutz-und-informationsfreiheit-in-der-pandemie-100.html>; Stand: 08.11.2021.
10. Medical Tribune. Bundesdatenschützer schickt Warnung an Krankenkassen: Elektronische Patientenakte nicht DSGVO-konform (12.11.2020) Im Internet: <https://www.medical-tribune.de/praxis-und-wirtschaft/praxismanagement/artikel/bundesdatenschuetzer-schickt-warnung-an-krankenkassen-elektronische-patientenakte-nicht-dsgvo-konfor/>; Stand: 08.11.2021.
11. Warnat-Herresthal S, Schultze H, Shastry KL et al. Swarm Learning for decentralized and confidential clinical machine learning. *Nature* 2021; 594: 265-270. DOI: 10.1038/s41586-021-03583-3
12. Rieke N, Hancox J, Li W et al. The future of digital health with federated learning. *npj Digital Medicine* 2020; 3: 119. DOI: 10.1038/s41746-020-00323-1
13. ZDF. Ethikrat: Weniger Datenschutz wäre vertretbar (29.10.2020) Im Internet: <https://www.zdf.de/nachrichten/digitales/coronavirus-warnapp-datenschutz--kritik-100.html> Stand: 16.06.2021.
14. Aggarwal R, Sounderajah V, Martin G et al. Diagnostic accuracy of deep learning in medical imaging: a systematic review and meta-analysis. *NPJ Digit Med* 2021; 4: 65. DOI: 10.1038/s41746-021-00438-z
15. Förch S, Klauschen F, Hufnagl P et al. Artificial Intelligence in Pathology. *Deutsches Arzteblatt international* 2021; 118: 194-204. DOI: 10.3238/arztebl.m2021.0011
16. Rush B, Celi LA, Stone DJ. Applying machine learning to continuously monitored physiological data. *J Clin Monit Comput* 2019; 33: 887-893. DOI: 10.1007/s10877-018-0219-z

17. Europäische Kommission. Lage der Union: Kommission schlägt neue Verordnung für das Gemeinsame Unternehmen für europäisches Hochleistungsrechnen vor – Fragen und Antworten (18.09.2020) Im Internet: https://ec.europa.eu/commission/presscorner/detail/de/qanda_20_1593; Stand: 29.08.2021.
18. Gemeinsame Wissenschaftskonferenz. Ausführungsvereinbarung zum GWK-Abkommen über die gemeinsame Förderung von Forschungsbauten, Großgeräten und des Nationalen Hochleistungsrechnens an Hochschule (26.11.2018) Im Internet: https://www.gwk-bonn.de/fileadmin/Redaktion/Dokumente/Papers/AV_FGH.pdf; Stand: 29.08.2021.
19. European Parliament and of the Council. Consolidated text: Regulation (EU) 2017/745 of the European Parliament and of the Council of 5 April 2017 on medical devices, amending Directive 2001/83/EC, Regulation (EC) No 178/2002 and Regulation (EC) No 1223/2009 and repealing Council Directives 90/385/EEC and 93/42/EEC (Text with EEA relevance) Im Internet: <https://eur-lex.europa.eu/legal-content/EN/ALL/?uri=CELEX:02017R0745-20200424>; Stand: 29.08.2021.

Abstract

The application of artificial intelligence (AI) is often associated with the use of large amounts of data for the construction of AI models and algorithms. This data should ideally comply with the FAIR Data principles, i.e. being findable, accessible, interoperable and reusable. However, the handling of health data poses a particular challenge in this context. In this article, we highlight the challenges of the data usage for AI in medicine using the example of Anaesthesia and Intensive Care Medicine. We discuss the current situation but also the obstacles for a wider application of AI in medicine in Europe and give suggestions how to solve the different issues. The article covers different subjects like data protection, research data infrastructures and approval of medical products. Finally, this article shows how it can nevertheless be possible to establish a secure and at the same time effective handling of data for use in AI at the European level despite its unneglectable difficulties.

Schlüsselwörter

Artificial intelligence, Research data, Data management, AI models, Data protection

CME-Fragen

1. Welche der genannten Datenbanken enthält multizentrische Daten?
 - a. MIMIC-III
 - b. eICU**
 - c. MIMIC-IV
 - d. HiRID
 - e. MIMIC-CXR
2. Welche Aussage zu öffentlich zugänglichen Datenbanken von Patientendaten ist korrekt.
 - a. Die Originalpublikationen der 3 Datenbanken wurden insgesamt bereits über 3600 Mal zitiert.

- b. Die Daten der eICU Datenbank haben eine deutlich höhere Auflösung als die der HiRID-Datenbank.
 - c. Die MIMIC-CXR Datenbank enthält lediglich radiologische Befunde, aber keine Bildaufnahmen.
 - d. **Die eICU-Datenbank enthält Datensätze zu einer größeren Anzahl Patienten als die HiRID-Datenbank.**
 - e. Alle im Text genannten Datenbanken enthalten lediglich US-amerikanische Daten.
- 3. Welche Antwort bezeichnet **kein** Prinzip der "FAIR Data principles"?
 - a. **Readability**
 - b. Findability
 - c. Interoperability
 - d. Reusability
 - e. Accessibility
- 4. Welche Aussage zum Datenschutz ist korrekt?
 - a. Die Verarbeitung von Gesundheitsdaten ist grundsätzlich erlaubt, es sei denn, für eine bestimmte Situation besteht ein explizites Verbot.
 - b. Die EU-Mitgliedstaaten können die Vorschriften der Datenschutz-Grundverordnung nicht verändern.
 - c. **Die meisten deutschen Bundesländer regeln die Verarbeitung von Gesundheitsdaten in ihren Krankenhäusern in speziellen Landesgesetzen.**
 - d. Gesundheitsdaten gehören nicht gemäß DS-GVO zu den „besonderen Kategorien von personenbezogenen Daten“.
 - e. An die Einwilligung von Patienten in die Verarbeitung ihrer Gesundheitsdaten werden keine besonderen Ansprüche gestellt.
- 5. Welche Aussagen zu den Methoden der KI trifft **nicht** zu?
 - a. Die für das Modell zur Verfügung stehenden Daten werden in mindestens einen Trainings- und einen Testdatensatz aufgeteilt.
 - b. **Ein Muster, welches in den Daten vorliegt, muss vor Training des Modells bereits genau bekannt sein.**
 - c. Die Qualität der Daten des Trainingsdatensatzes ist für das resultierende Modell von entscheidender Bedeutung.
 - d. Existierende Formeln zu physiologischen Prozessen haben häufig eine gleich gute bis bessere prädiktive Fähigkeit als Machine Learning Modelle.
 - e. Das Föderale Maschinelle Lernen stellt eine Technik dar, mit der viele Probleme mit dem Datenschutz umgangen werden können.
- 6. Welche Aussage zum Digital Object Identifier (DOI) ist **nicht** korrekt?
 - a. Durch einen DOI kann ein Objekt eindeutig identifiziert und gefunden werden, auch wenn sich sein physischer Standort ändert.
 - b. Der DOI baut auf dem Handle-System auf.
 - c. Ein Datensatz, der über den Service B2SHARE gespeichert wird, erhält automatisch einen DOI.
 - d. Die Vergabe eines DOI dient auch der Gewährleistung der FAIR Data Prinzipien.
 - e. **Ein DOI kann lediglich veröffentlichten Manuskripten zugewiesen werden.**
- 7. Welche Aussage zu Hochleistungsrechnen (high performance computing, HPC) ist korrekt?
 - a. Hochleistungsrechner weisen alle einen einheitlichen Aufbau auf.
 - b. Der Ausbau von Kapazitäten des HPC hat in der EU in den letzten Jahren in der Politik praktisch keine Rolle gespielt.

- c. **Typische Anwendungsbereiche für HPC in der Medizin sind die Bildverarbeitung und die Analyse kontinuierlich erhobener Vitalparameter.**
 - d. Die Bündelung von HPC-Kapazitäten wird in Europa durch das Gauss Centre for Supercomputing koordiniert.
 - e. Hochleistungsrechnen wurde in der Medizin von Anfang an in großem Umfang eingesetzt.
8. Welche Aussage zur European Science Cloud (EOSC) ist **nicht** korrekt?
- a. **Die Daten, die innerhalb des EOSC Projektes verwendet werden, entsprechen nicht den FAIR Data Prinzipien.**
 - b. Ziel der EOSC ist es, 72 bestehende Forschungsdateninfrastrukturen in Europa zu vereinigen
 - c. EOSC LIFE stellt eine besondere Spezialisierung der EOSC für die Life sciences dar.
 - d. Die EOSC ist ein Projekt der Europäischen Union.
 - e. Weitere Datenmanagementdienste wie die EUDAT Collaborative Data Infrastructure sind an das Projekt integriert.
9. Welche Aussage zu Software als Medizinprodukt ist richtig?
- a. Die Einstufung von Software als Medizinprodukt in die entsprechenden Risikoklassen ist seit 2021 im Medizinproduktegesetz (MPG) geregelt.
 - b. Einer Software-Anwendung, die als Medizinprodukt genutzt werden soll, gehört immer zur Risikoklasse I.
 - c. Software, deren Einsatz im schlimmsten Fall den Tod eines Patienten bedingen kann, darf in der EU nicht zugelassen werden.
 - d. Software, deren Anwendung einen chirurgischen Eingriff notwendig machen kann, gehört zur Risikoklasse IIa.
 - e. **Software, die Informationen liefert, die Ärzte bei Entscheidungen zur Diagnose oder Therapie eines Patienten unterstützen sollen, gehört im Allgemeinen zur Risikoklasse IIa.**
10. Welche Aussage zur Zulassung von Medizinprodukten ist **nicht** richtig?
- a. Die Möglichkeit, Medizinprodukte in erster Linie auf der Basis von Ergebnissen aus der Literatur zuzulassen, wurde in der MDR deutlich eingeschränkt.
 - b. **Die Möglichkeit der Zulassung eines Medizinproduktes nach dem Äquivalenzprinzip ist mit Inkrafttreten der MDR abgeschafft worden.**
 - c. Bei der Zertifizierung von Medizinprodukten der Klasse III gilt eine klinische Prüfung praktisch als Regelfall.
 - d. Software, die Ärzte bei der Therapie von Patienten berät, muss ein Konformitätsverfahren bei einer Benannten Stelle durchlaufen.
 - e. Die MDR verpflichtet die Hersteller eines Medizinproduktes auch zur Überwachung des Produktes nach der Markteinführung.