

KI-Solutionismus.

Erarbeitet von
Titel Marie von Lobenstein

Das Thema „KI-Solutionismus“ ist äußerst komplex und wird aus verschiedenen Fachdisziplinen und Perspektiven betrachtet. In diesem Video haben wir daher eine Auswahl getroffen, um einen ersten Überblick zu geben. Eine vollständige Behandlung des Themas würde den Rahmen der Lehrveranstaltung sprengen und ist im Rahmen dieses Kurses nicht möglich.

Lernziele	1
Inhalt	2
Einstieg	2
Die Fallstricke des KI-Solutionismus	3
Ein ausgewogener KI-Lösungsansatz	4
Take-Home Message	5
Quellen	5
Weiterführendes Material	5
Disclaimer	5

Lernziele

- Du kannst KI-Solutionismus definieren
- Du kannst Beispiele für Versprechen und Fallstricke von KI-Solutionismus nennen
- Du kannst erklären, warum KI nicht immer die „bessere“ Lösung ist für ein gesellschaftliches Problem

Inhalt

Einstieg

Wenn man sich umhört, ist im Bereich Nachhaltigkeit die Hoffnung groß, dass KI uns helfen kann, unsere Erde zu schützen und beispielsweise den Klimawandel zu stoppen oder Ungleichheiten zu beseitigen. Gleichzeitig ist schon das Programmieren und Nutzen von künstlicher Intelligenz mit einem erheblichen Ressourcen- und Energieverbrauch verbunden und wir wissen auch, dass KI diskriminieren kann.

Ich möchte in diesem Video mit euch über die Annahme oder das Versprechen reden, dass KI die Lösung für all unsere menschengemachten Probleme bereithält. Dabei wollen wir nicht pessimistisch werden in Bezug auf Tech-Lösungen. Wir wollen uns kritisch mit der Annahme auseinandersetzen, dass wir durch genügend Daten und Machine Learning Algorithmen ein Allheilmittel für alle Probleme unserer Gesellschaft gefunden haben. Diese Annahme wird auch als KI-Solutionismus bezeichnet.

Evgeny Morozov ist einer der Kritiker*innen von Tech-Solutionismus und hat den Begriff stark geprägt.

Das Problem an Tech-Solutionismus ist für ihn, dass vorausgesetzt wird, dass Technologien komplexe, soziale Probleme – von der Gesundheitsfürsorge, über Gleichberechtigung bis hin zum Klimawandel – mit relativer Leichtigkeit durch berechenbare Lösungen optimieren oder beheben können, ohne dabei die Probleme als Solche zu hinterfragen. So einfach und ohne Nebeneffekte sind diese „einfachen“ Lösungen in der Realität laut ihm nämlich meist nicht. Für Morozov geht es darum, Probleme genauer zu betrachten und nicht nur das System, indem sie auftauchen, zu optimieren.

Quelle [1][2]

Wenn wir Morozovs Tech-Solutionismus auf KI anwenden reden wir also von dem Glauben komplexe Lösungsansätze durch Automatisierung zu vereinfachen, indem man einzelne (menschliche) Schritte im Prozess ‚überspringt‘, die aber entscheidend für einen fairen, verantwortungsvollen und nachhaltigen Umgang mit der Technologie sein könnten.

Obwohl KI in verschiedenen Bereichen zweifellos bedeutende Fortschritte gemacht hat, ist es wichtig, die Annahme, dass KI die Lösung beispielsweise für unsere Nachhaltigkeitsziele ist, umfangreich zu betrachten.

Lasst uns beginnen, indem wir uns anschauen, was KI-Solutionismus eigentlich unter anderem verspricht.

Grob gesagt verspricht KI-Solutionismus, dass wir bedeutende Fortschritte bei der Linderung komplexer Dilemmata erzielen können, wenn nicht sogar, sie vollständig zu beheben, indem wir ihre Kernprobleme auf einfachere ingenieurtechnische Probleme reduzieren.

Diese Aussicht auf eine simple Lösung durch KI ist für uns natürlich verlockend, weil diese Kernprobleme so in unserer sonst so komplizierten Welt, einfach und geradlinig bewältigt werden können. Außerdem wird durch KI-Ansätze oft eine erschwingliche Lösung versprochen, in einer Welt mit knappen Ressourcen und limitierten finanziellen Möglichkeiten. Eine dritte zu bedenkende Perspektive ist, dass KI-Solutionismus unseren Optimismus bezüglich Innovation verstärkt.

Aber was ist denn eigentlich problematisch daran, KI als Lösung unsere Probleme zu sehen oder zu verkaufen? Ich zeige euch hier mal ein paar Beispiele für Fallstricke des KI-Solutionismus.

Die Fallstricke des KI-Solutionismus

1. Überhöhte Erwartungen: Der Hype um KI kann zu unrealistischen Erwartungen führen. Wenn KI keine wundersamen Ergebnisse liefert, kann dies das Vertrauen untergraben und zu Desillusionierung führen. Außerdem kann dieser Hype auch blenden. Nicht jedes Problem kann durch eine „simple“ KI-Lösung auch wirklich gelöst werden. In manchen Fällen trägt KI auch weiter zu den bestehenden Problemen bei.
2. Falsche Versprechen: Evgeny Morozov beschreibt ein weiteres Problem: Solutionismus hört nicht tatsächlich auf diejenigen, deren Probleme es vorgibt zu behandeln, sondern bevorzugt es, ihnen lediglich vorgefertigte Lösungen zu liefern. Nehmen wir das Bildungswesen zum Beispiel. Durch KI sollen viele nervige Prozesse optimiert werden und eine individualisierte Bildung ermöglicht werden. Zum Beispiel kann KI als Korrekturhilfe genutzt werden, um Lehrende zu unterstützen und ihnen mehr Zeit für andere Aufgaben zu ermöglichen. Dieses Korrekturtool mag sogar eine Zeitersparnis sein, darf aber nicht dazu führen, dass die Einwände von Lehrenden in Bezug auf Lehrqualität, Ressourcen oder Reformen ignoriert werden, weil es ja vermeintliche KI-Lösungen gibt, die auf diese Einwände bereits eingehen.
3. Ein Problem, wo kein Problem ist: Der Solutionismus geht davon aus, dass es überhaupt Probleme gibt, die gelöst werden müssen. Das ist ein ganz spezifisches (ingenieursgeleitetes) Denkmodell, das nicht unbedingt auf andere Bereiche des Lebens übertragbar ist (z. B. sind viele kulturelle, philosophische, ästhetische etc. Bereiche gar nicht unbedingt durch das Modell ‚Problem‘ – ‚Lösung‘ zu fassen, sondern erfordern andere Zugänge).

4. KI ist Teil des Problems: Viele KI-Algorithmen werden auf der Grundlage historischer Daten trainiert, die Verzerrungen enthalten können. Dies kann bestehende soziale, rassistische oder geschlechtsspezifische Vorurteile aufrechterhalten und sogar noch verschärfen, was zu ungerechten Ergebnissen führt. KI verbraucht viel Energie und natürliche Ressourcen. Also anders gesagt, KI ist nicht nur eine mögliche Lösung, sondern, schon Teil des Problems. Außerdem erfordern KI-Lösungen häufig laufende Wartungen, Updates und erhebliche Rechenressourcen. Die Implementierung und Aufrechterhaltung von KI-Systemen können kostspielig und komplex sein.
5. Wer profitiert von AI-Solutionismus : Die Big Player der KI-Branche freuen sich natürlich, wenn der Eindruck verstärkt wird, dass KI-Ansätze zur Problemlösung als effektiver gelten, als Alternativen mit sozialen und politischen Dimensionen. Mit KI-Tech wird viel Geld gemacht. Das Versprechen: Für jedes Problem gibt es eine KI-basierte Lösung oder ein Anwendungsabo, welches man abschließen kann. Man sollte sich also immer fragen, wer profitiert davon, dass ich KI nutze und gibt es nicht eventuell auch schon eine andere Lösung oder einen Lösungsansatz, die man in Betracht ziehen kann.

Wie so oft, ist es auch hier ratsam sich irgendwo zwischen den Versprechen die KI-Solutionismus macht und den Problemen, die es aufwirft, zu treffen. Wir wollen keine Tech-Pessimist*innen werden, aber neugierig, differenziert und wachsam bleiben.

Wie können wir zum Beispiel über eine Zukunft mit KI nachdenken?

Ein ausgewogener KI-Lösungsansatz

Der Schlüssel zur Nutzung des Potenzials der KI bei gleichzeitiger Abschwächung ihrer Fallstricke liegt in einem ausgewogenen Ansatz, dafür hier ein paar Vorschläge:

1. Realistische Erwartungen: Es ist wichtig, KI als ein Werkzeug zu sehen, das die menschlichen Fähigkeiten ergänzt, und nicht als magisches Allheilmittel. Es ist wichtig, Grenzen und Fähigkeiten zu verstehen
2. Ethische Erwägungen: Entwickler*innen und Unternehmen müssen der ethischen Entwicklung von KI Priorität einräumen und sich mit Fragen der Voreingenommenheit, Fairness und Transparenz auseinandersetzen. Zu einer verantwortungsvollen KI-Entwicklung gehören vielfältige Teams und eine ständige Überprüfung.
3. Regulierung und Steuerung: Die politischen Entscheidungsträger*innen sollten klare Regelungen und Governance-Rahmen für KI-Technologien schaffen, um die Verantwortlichkeit zu gewährleisten und die Rechte des Einzelnen zu schützen. Dabei müssen die, die von den Lösungen profitieren sollen, in die Entscheidungsfindungen mit eingebunden werden.

Take-Home Message

Zusammenfassend lässt sich sagen, dass der KI-Lösungsansatz vielversprechend ist, aber auch erhebliche Risiken birgt. Es ist unerlässlich, sich KI mit einer differenzierten Perspektive zu nähern und sowohl die potenziellen Vorteile als auch die Grenzen anzuerkennen. Ein ausgewogenes Verhältnis zwischen technologischem Fortschritt und verantwortungsvollem Einsatz ist der Schlüssel zur Ausschöpfung des wahren Potenzials von künstlicher Intelligenz in einer sich rasch entwickelnden Welt. Man sollte sich immer, wenn KI als ‚schnelle/ einfache‘ Lösung für ein Problem dargestellt wird, fragen, ob es sich hier überhaupt um ein klassisches Problem handelt, ob das Versprechen eines ‚quick fix‘ diesem angemessen ist und welche Interessen möglicherweise hinter einer Lösungsfindung stecken könnten.

Quellen

- Quelle [1] Morozov, E. (2013). *To save everything, click here: The folly of technological solutionism*. PublicAffairs.
- Quelle [2] Münch, M., & Morozov, E. (2013, October 23). Interview Evgeny Morozov: "Uns steht eine Datenapokalypse bevor." *Bundeszentrale für politische Bildung*. <https://www.bpb.de/themen/medien-journalismus/netzdebatte/171007/interview-evgeny-morozov-uns-steht-eine-datenapokalypse-bevor/>.

Weiterführendes Material

Lindgren, S., & Dignum, V. (2023). Beyond AI solutionism: toward a multi-disciplinary approach to artificial intelligence in society. In *Handbook of Critical Studies of Artificial Intelligence* (pp. 163-172). Edward Elgar Publishing.

Deutschlandfunkkultur.de. (n.d.). *Philosophie des Silicon Valley - der Geist des Digitalen Kapitalismus*. Deutschlandfunk Kultur. <https://www.deutschlandfunkkultur.de/philosophie-des-silicon-valley-der-geist-des-digitalen-100.html>

Disclaimer

Transkript zu dem Video „KI und Nachhaltigkeit: KI-Solutionismus“, Marie von Lobenstein. Dieses Transkript wurde im Rahmen des Projekts ai4all des Heine Center for Artificial Intelligence and Data Science (HeiCAD) an der Heinrich-Heine-Universität Düsseldorf unter der Creative Commons Lizenz [CC-BY 4.0](https://creativecommons.org/licenses/by/4.0/) veröffentlicht. Ausgenommen von der Lizenz sind die verwendeten Logos, alle in den Quellen ausgewiesenen Fremdmaterialien sowie alle als Quellen gekennzeichneten Elemente.