

KI-Kompetenz: Mythen und Fakten über ChatGPT

Transkript zum Video

Zwischenüberschriften und Sprecherbezeichnungen wurden zur besseren Orientierung redaktionell ergänzt.

Dieses Transkript wurde mithilfe von KI erstellt und kann vereinzelt Ungenauigkeiten enthalten.

00:00 – Begrüßung und Einführung

Simon Berlin

Hallo und herzlich willkommen zu einem kurzen Crashkurs, der hoffentlich ein kleines bisschen KI-Kompetenz vermittelt. Mein Name ist Simon Berlin. Ich bin Autor für die Süddeutsche Zeitung und schreibe einen Newsletter, der „Social Media Watchblog“ heißt. Und in beiden Publikationen versuche ich das Internet so zu erklären, dass es meine Eltern verstehen. Das heißt, in den vergangenen drei Jahren ist das zunehmend von KI, künstlicher Intelligenz, geprägt, weil mir da ganz viele Fragen begegnen. Sowohl in Gesprächen mit Leserinnen und Freunden fallen mir immer wieder bestimmte Mythen und Missverständnisse auf. Und mit einigen davon würde ich heute gerne aufräumen und sie ansprechen. Dafür habe ich einen kurzen Vortrag vorbereitet, der so, ich würde sagen, fünfzehn bis zwanzig Minuten dauern wird und danach bleibt noch genug Raum für Fragen. Ich starte jetzt mal die Präsentation.

Das dauert jetzt ganz kurz. Nein, erst muss ich meinen Bildschirm teilen.

[Stille]

Und jetzt sollte die erste Folie zu sehen sein. Damit fange ich an. Das ist jetzt das Thema, das ich gerade schon vorgestellt hatte. Ich fange mit dem quasi ersten, wenn man so will – das ist kein Mythos, aber eine Sache, die mir auffällt, die mich am Diskurs über KI stört, an.

01:29 – KI ist ein Oberbegriff

Simon Berlin

Und das ist schon der Begriff KI, künstliche Intelligenz, weil das ist ein fürchterlich schwammiger und letztlich eigentlich nichts sagender Begriff, der seit den Sechzigern existiert und jetzt so pauschal für irgendwie alles verwendet wird und oft auch falsch genutzt wird. Besonders problematisch finde ich, wenn von „der KI“ oder „einer KI“ die Rede ist. Das gibt es in der Form einfach nicht. Medien sind da zum Teil auch mitschuldig, ich nehme mich da gar nicht aus, einfach weil es sich irgendwie so eingebürgert hat und so einfach ist und KI auch so ein schön kurzes Wort ist. Das passt so wunderbar in Überschriften. Ich habe vor Kurzem dazu die deutsche

Informatikerin Katharina Zweig interviewt, die mir was sehr Interessantes gesagt hat, was ich hier kurz zitieren möchte. Sie sagte: „KI ist kein Ding, sondern eine Forschungsrichtung, eine Menge unterschiedlicher Methoden.

Die Systeme, die diese Methoden verwenden, nennen wir leider auch oft KI. Das ist irreführend. Die KI gibt es nicht. Es ist ein Computer oder ein System, eine Maschine oder eine Software.“ Man sieht hier, dass quasi KI oder jetzt AI auf dieser Folie ein Oberbegriff ist, der für sehr viele unterschiedliche Dinge steht. Und das muss man sich jetzt nicht alles merken, diese ganzen Fachbegriffe. Was wichtig ist, ist eben sich zu merken: KI ist ein Oberbegriff. Und wenn wir heute über KI reden, dann ist in den allermeisten Fällen generative KI gemeint, die wiederum meist auf Sprachmodellen oder auch LLMs beruht. Das sind die Sachen, die man sich merken sollte.

03:01 – Chatbots sind keine verlässlichen Quellen

Simon Berlin

Damit zum zweiten Punkt. Immer wieder in Unterhaltungen oder auch im Internet in Foren lese ich so Sachen, oder höre ich so Sachen wie: „Ja, ich habe mal KI gefragt“, oder „aber ChatGPT sagt doch.“ Und das ist was, was mir wirklich auf die Nerven geht. Ich finde, dass Chatbots nie als Quelle für irgendwas herhalten sollten, weil diese Sprachmodelle am laufenden Band konfabulieren. Also sie füllen Wissenslücken durch Erfindungen und sie sind einfach nicht zuverlässig. Der amerikanische Tech-Analyst Benedict Evans hat das ziemlich gut auf den Punkt gebracht. Der schreibt nämlich: „Ich brauche keine Antwort, die wahrscheinlich richtig ist. Ich brauche eine Antwort, die garantiert richtig ist. Diese Modelle liefern keine richtigen Antworten.

Sie sind probabilistische, statistische Systeme, die berechnen, wie eine plausible Antwort aussehen könnte.“ Deshalb werden Sprachmodelle häufig in der Forschung auch „stochastische Papageien“ genannt, die einfach nur nachplappern. Das trifft es ganz gut. Allerdings finde ich eine wichtige Einschränkung noch: Das macht Sprachmodelle überhaupt nicht unbrauchbar. Es bleibt eine sehr, sehr mächtige Technologie, aber – und das ist eben wichtig, man muss wissen, wofür man sie besser nicht einsetzt. Und immer dann, wenn es um, quasi um Fakten geht, dann sollte man sich nicht allein auf KI verlassen.

04:29 – Sprachmodelle sind komplex

Simon Berlin

Der dritte Punkt, den ich ansprechen möchte, ist ein sehr, wirklich weit verbreitetes Missverständnis, und zwar, dass ChatGPT oder andere Chatbots einfach sind. Ich verstehe, woher das kommt, weil man tippt oder spricht irgendeinen Text ein und die Maschine antwortet. Ja, was soll da schon schiefgehen? Das versteht ja wirklich jeder. Das Dumme ist, jede Menge kann da schiefgehen. Tatsächlich sind eben Sprachmodelle oder Chatbots eigentlich hochkomplex in der Bedienung und im Verständnis. Simon Willison, ein britischer Programmierer, der sich sehr

intensiv mit KI beschäftigt, sagt dazu: „Sprachmodelle sind Kettensägen, die als Küchenmesser daherkommen. Sie wirken einfach, aber man braucht Wissen und Erfahrung, um sie zu bedienen.“ Ich möchte noch jemanden zitieren, einen amerikanischen Professor, der zu KI forscht und lehrt, Ethan Mollick.

Der sagt: „Viele Menschen spielen mit KI herum, aber wissen nicht, was sie damit anfangen sollen. Meine Erfahrung zeigt, man muss etwa zehn Stunden mit einem KI-System verbringen. Erst dann entwickelt man ein Gefühl dafür, wobei einem KI im Leben und Beruf helfen kann.“ Und das deckt sich sehr genau mit meiner Erfahrung, sowohl mit meiner Erfahrung im persönlichen Einsatz, wie ich KI nutze, aber auch in meiner Erfahrung, wie ich quasi andere Menschen beim Einsatz von KI oder halt eben Chatbot-Sprachmodellen beobachte. Dass man quasi in der Anfangszeit oft gar nicht so richtig weiß, was das eigentlich soll, weil man einfach Zeit benötigt, um das herauszufinden. Und dann kommt noch dazu eine Sache, dass sich Sprachmodelle wahnsinnig schnell verändern. Das macht die Bedienung noch schwerer, einfach weil die Technologie sich gerade so schnell entwickelt.

Das heißt, vermeintliche Gewissheiten, die vor einem halben Jahr oder vor einem Jahr gegolten haben, sind vielleicht jetzt gar nicht mehr aktuell. Also man muss ständig auf dem aktuellen Stand bleiben. Allein so Dinge, die total banal wirken, wie: „Welches Modell wähle ich eigentlich aus?“ Oder: „Wie prompte ich es?“ Also was schreibe ich quasi da rein, sind super komplex und ändern sich ständig. Deshalb als Take-away: Nicht von so einem vermeintlich simplen Design täuschen lassen. KI ist wirklich komplex.

06:49 – Macht KI dumm?

Simon Berlin

Der vierte Punkt, den ich gerne ansprechen möchte, ist etwas, was mir häufig so im vergangenen Jahr in Form von Überschriften in Medien entgegengekommen ist, was viele Schlagzeilen gemacht hat. Und das ist die Annahme, dass der Einsatz von KI dumm macht und in irgendeiner Form, ja, die Intelligenz verkümmern lässt. Wie so oft, es hat einen wahren Kern. Tatsächlich kann der Einsatz von KI, wenn man sich komplett darauf verlässt, bestimmte kognitive Fähigkeiten oder auch so das kritische Denken verkümmern lassen. Da gibt es diverse Studien dazu, die das gezeigt haben. Eine davon, also eine Autorin dieser Studien, ist hier Nataliya Kosmyna. Die hat am sehr renommierten MIT in Boston an der Studie mitgeschrieben, die hieß „Your Brain on ChatGPT“.

Und die haben dabei Probandinnen und Probanden unter anderem in einen Hirnscanner geschickt und geschaut, was passiert, wenn man bestimmte Aufgaben mit KI und ohne KI durchführt. Und dann kam eben raus, dass beim Einsatz von KI weniger Hirnareale aktiv sind und bestimmte Zentren gar nicht aktiviert werden. Was dann wiederum eben sehr viele Medien zu den Schlagzeilen veranlasst hat, dass der Einsatz von ChatGPT Menschen dümmer macht. Und dann hat aber quasi extra eben, haben die Autor:innen dieser Studie dann noch mal nachträglich quasi darauf reagiert und auch ihrer Studie noch mal so eine Art Handreichung hinzugestellt, wo sie auf, ja, viele Fragen und so Fehlannahmen eingehen. Unter anderem hat sie dann gesagt, auf die Frage: Kann man sagen, dass Sprachmodelle uns dümmer machen? „Nein!“

Bitte verwenden Sie keine Wörter wie „dumm“, „Schaden“, „Passivität“, „geistiger Verfall“ oder Ähnliches. Diese Begriffe haben wir in der Studie bewusst nicht verwendet. Das gilt insbesondere, wenn Sie als Journalist darüber berichten.“ Ich habe das selbst gemerkt, wie verführerisch das ist, wenn man so eine Studie liest, weil man dann sofort denkt: „Ah ja, okay, krass, so eine Schlagzeile würde ja sicher funktionieren. Die interessiert ja sicher viele Leute.“ Aber eben diese und viele andere Studien geben das halt gar nicht her. Die Studie sagt erst mal nur aus: Bestimmte Hirnareale sind nicht aktiv, wenn man KI nutzt. Punkt. Das ist keine Aussage darüber, dass KI uns dümmer macht. Was allerdings, das hatte ich ja schon anfangs gesagt, sehr wohl der Fall ist, dass wenn man sich nur auf KI verlässt, möglicherweise die eigene so Problemlösungskompetenz und die Kreativität darunter leiden.

Also schließt sich irgendwie die Frage an: Wie kann man KI gewinnbringend einsetzen, so, dass es vielleicht sogar das eigene Denken bereichert? Und da hat die deutsche Mathematikerin und KI-Dozentin Barbara Lampl was gesagt, was ich eigentlich in der Kürze sehr gut auf den Punkt gebracht finde. Und zwar sagt sie: „Hör auf zu denken, dass KI es für mich macht. Fang an zu denken, dass KI mir hilft, es besser zu machen.“ Das finde ich einen, einen sehr pragmatischen und sinnvollen Ansatz, der sich auch total mit allen mir bekannten Studien und diversen Metaanalysen zu diesem Thema deckt. Es kommt nämlich nicht darauf an, ob man KI nutzt, sondern es kommt darauf an, wie man KI nutzt. Und KI oder Sprachmodelle an sich schaden unserem Gehirn nicht, aber komplett gedankenloser Gebrauch kann eben unser Denken beeinträchtigen.

Also das Denken nicht an ChatGPT auslagern, sondern immer versuchen, selbst zu denken und Sprachmodelle nur als Unterstützung zu verwenden.

10:46 – KI-Blase und wirtschaftliche Abhängigkeiten

Simon Berlin

Zum nächsten Punkt, Punkt fünf. Auch was, was jetzt gerade, aber auch so seit, seit einem halben Jahr viel durch Medien geht und das ist dieser, der Diskurs über die Blase, die KI angeblich ist. Da werden viele Vergleiche zur Dotcom-Blase um die Jahrtausendwende gezogen, und die ja damals, die Älteren werden sich erinnern, dann doch ziemlich heftig geplatzt ist, was durchaus krasse Folgen für die Weltwirtschaft hatte. Diese Wortwahl „Blase“ suggeriert ja so ein bisschen, da ist nur Luft drin, das sei eine Technologie ohne Substanz. Mittlerweile ist es sogar so weit, dass Sam Altman, also der Chef von OpenAI, höchstpersönlich, das ist ein wörtliches Zitat von ihm, sagt: „Ja, natürlich ist KI eine Blase. Start-ups, die aus drei Leuten mit einer Idee bestehen, sammeln dreistellige Millionensummen ein. Das ist verrückt.

Jemand wird sich die Finger verbrennen.“ Und ich stimme nicht oft mit Sam Altman überein, aber hier stimme ich ihm total zu. Also er hat da definitiv einen Punkt. Wenn man noch einen Schritt weitergeht: Das sind jetzt hier zwei Abbildungen aus zwei renommierten US-Medien, dem Wall Street Journal und Bloomberg. Keine Sorge, das muss niemand im Detail verstehen. Ich verstehe es auch nur mit großer Mühe, aber da sind so die, die Details der aktuellen KI-Deals mal skizziert

in immer noch vereinfachter Form. Und man sieht, da fließen sehr viele Linien von links nach rechts und im Kreis, was einfach deutlich macht, wie eng diese ganzen Unternehmen miteinander verwoben sind, wer da in wen investiert. Und wenn man dann auch weiß, dass es hier nicht um ein paar Millionen geht, sondern in der Summe um mehrere Billionen – und das ist kein Übersetzungsfehler.

Also wir sprechen nicht über Milliarden, sondern wir sprechen wirklich über Billionen, die in den kommenden Jahren gerade in KI-Infrastrukturen, Rechenzentren fließen, dann ist es schon eine krasse wirtschaftliche Abhängigkeit mit hochkomplexen Verträgen, die auch für die gesamte Wirtschaft hoch riskant ist. Also das kann man sich schon merken: Sehr viel von der gerade US-amerikanischen Realwirtschaft beruht aktuell auf Investitionen in KI. Aber, und darauf wollte ich eigentlich hinaus, das macht KI als Technologie nicht wertlos. Also selbst wenn ein Teil dieser Investitionen Blasencharakter hat, dann bleibt trotzdem im Zentrum eine Technologie, die Substanz hat.

Wenn man sich in der Wirtschaftsgeschichte vergangene Blasen anschaut, zum Beispiel die Eisenbahn, wo es auch blasenartige Investitionen gab oder die Dotcom-Blase, dann haben am Ende die Technologien, die im Zentrum standen, trotzdem die Gesellschaft verändert. Also von der Eisenbahn ist ja heute noch was übrig und das war wichtig. Und die Dotcom-Blase hat zum Beispiel dazu geführt, dass massenweise Glasfasernetze ausgebaut wurden, die überhaupt dazu geführt haben, dass das Internet von heute so existiert, wie es eben ist. Also nur weil etwas eine Blase ist oder teilweise auch Spekulanten investieren, entwertet es die Technologie nicht.

13:59 – Umweltwirkungen von Chatbots

Simon Berlin

Der vorletzte Punkt. Auch was, was mir sehr häufig begegnet. Das ist die Annahme, dass der Einsatz von Chatbots der Umwelt schadet. Tatsächlich kenne ich einige Menschen, die aktiv ChatGPT meiden, weil sie Angst haben, sehr viel Strom und Wasser zu verbrauchen. Ich habe das tatsächlich längere Zeit auch selbst für ein Problem gehalten und dachte: „Oh Gott, meine Klimabilanz ist vermutlich eh schon fürchterlich. Jetzt füge ich mit dem Einsatz von AI noch was hinzu.“ Dann habe ich angefangen, mich intensiver mit den Studien zu beschäftigen dazu.

Und jemand, der da wirklich tief reingegangen ist und irgendwie alles ausgewertet hat, das ist ein britischer Physiker – nein, ich glaube, der ist gar nicht Brite, der ist Amerikaner, ein amerikanischer Physiker, Andy Masley, der in unfassbar langen Artikeln das alles dargelegt hat und die eine Sache, die er zitiert, die es, finde ich, gut auf den Punkt bringt. Er sagt: „Sie können ChatGPT so viel nutzen, wie Sie möchten, ohne sich Sorgen zu machen, dass Sie dem Planeten schaden. Sich über die persönliche Nutzung von ChatGPT Gedanken zu machen, ist verschwendete Zeit, die sie stattdessen den drängenden Problemen des Klimawandels widmen könnten.“ Es ist wirklich so, quasi einfach nur ChatGPT oder ein anderes Sprachmodell nach irgendwas zu fragen, verbraucht verschwindend geringe Mengen an Strom und Wasser. Das ist nicht das Problem.

Es gibt tausend andere Verhaltensweisen, die viel, viel, viel schädlicher sind. Es gibt eine Einschränkung, die man noch machen kann und sollte: Wenn man massenhaft längere KI-generierte Videos erzeugt, dann ist das tatsächlich relativ ressourcenintensiv, aber nachdem es die allermeisten Menschen nicht tun, kann man wirklich so für die persönliche individuelle Nutzung mitnehmen: Das ist nicht das Problem. Das ist nicht die Verhaltensänderung, die man anstreben sollte. Eine Anmerkung noch dazu: Als quasi Gesamtheit, wenn man die ganze Infrastruktur anschaut, die für den Betrieb von Sprachmodellen nötig ist, also diese gewaltigen Rechenzentren und das teilweise monatelange Training, dann ist das schon ein Faktor, der ins Gewicht fällt, allerdings im Vergleich zu anderen Industrien auch nur ein sehr, sehr kleiner Faktor.

16:27 – Reale Risiken statt Science-Fiction

Simon Berlin

Und damit komme ich zum letzten Punkt, der seit Jahren sich leider, würde ich sagen, aus meiner Perspektive hält im Diskurs. Und das ist die Annahme, dass eine der größten Gefahren, die von generativer KI ausgehen, in intelligenten Maschinen besteht. Es wird seit Jahren gewarnt vor sogenannter AGI, also Artificial General Intelligence. Das sind Maschinen, die der menschlichen Intelligenz überlegen sind. Und häufig werden dann so Terminator-Szenarien, also Science-Fiction-Szenarien zitiert, wo Maschinen dann die Macht übernehmen und Menschen unterdrücken. Ich halte das für eine ziemlich strategisch lancierte Scheindebatte, die vor allem in der Anfangszeit, so nach der Veröffentlichung von ChatGPT, Ende 2022, im Jahr 2023, ein großartiges PR-Instrument war. Man könnte nämlich sagen: „Schaut mal, wie mächtig unsere Technologie ist.“

Gebt uns unbedingt mehr Geld, weil ihr ja an Unternehmen wie OpenAI beteiligt sein wollt, wenn diese Technologie so mächtig ist, dass sie im schlimmsten Fall auch solche Konsequenzen haben könnte. Aber vertraut uns, wir tun alles dafür, dass sie sicher ist.“ Sam Altman war eben derjenige, der das am häufigsten und lautstärksten gesagt hat. Mittlerweile ist es so ein bisschen in den Hintergrund getreten, aber es bleibt trotzdem noch ein Narrativ, was immer wieder erwähnt wird. Das größte Problem, was ich damit habe, ist, dass es von den sehr realen Problemen und Risiken von KI ablenkt, die auch überhaupt nicht hypothetisch sind, sondern die jetzt schon real sind. Und da möchte ich noch kurz auf drei der sehr vielen Probleme eingehen, die ich gerade bei KI sehe, die mir deutlich mehr Sorgen machen als Terminator. Das hier ist das Erste, was mich gerade beruflich sehr beschäftigt.

Das sind jetzt vier Personen, die man hier sieht, die sich alle das Leben genommen haben, nachdem sie sich längere Zeit mit ChatGPT oder anderen KI-Systemen unterhalten haben. Und man muss zwar dazu sagen, nein, ChatGPT oder OpenAI ist nicht schuld daran. Es war sicher nicht ursächlich, aber ich glaube schon, in allen vier Fällen, die ich relativ genau kenne, kann man dazu sagen, dass diese Sprachmodelle zumindest dazu beigetragen haben, dass die Menschen Suizid begangen haben. Also was man – das ist jetzt wirklich nur die Spitze des Eisbergs – ich glaube, man kann definitiv festhalten, dass viele Menschen noch nicht gelernt haben, mit Maschinen zu kommunizieren, die so tun, als seien sie Menschen.

Und das führt eben zu meiner Schlussfolgerung, dass generative KI ein wirklich gigantisches Sozialexperiment mit völlig ungewissem Ausgang ist und wir, glaube ich, gerade erst beginnen, die Folgen zu sehen. Der zweite Punkt ist das Stichwort Desinformation. Dieses Bild kennt vielleicht die eine oder der andere. Das war ein Meme, was vor ein paar Jahren mal um die Welt ging. Das ist der damalige Papst im weißen Balenciaga-Daunenparka und es war ein Fake. Aber das war so das erste bekannte Beispiel, wo ein Fake viral ging, ohne dass sofort alle erkannt haben, dass es ein Fake ist. Und damals war es sowohl noch irgendwie ein bisschen lustig, weil auch irgendwie harmlos und es war vor allem relativ aufwendig, das Bild zu machen. Mittlerweile, ich habe da so ein bisschen rumgespielt, einfach zwei Fotos von mir, weil ich niemand anderen faken wollte.

Und zwei Minuten später habe ich einen weißen Daunenparka an und es sieht wirklich gut aus. Sogar der Schattenwurf ist gut und das ist jetzt totaler Low Tech. Das könnte jeder. Das ist einfach nur mit so einem aktuellen Bildmodell von Google rumgespielt. Also es ist wahnsinnig einfach geworden, nicht nur gefälschte Bilder, sondern mittlerweile auch gefälschte Videos zu veröffentlichen. Und auch das wird noch krasse Herausforderungen mit sich bringen, weil einfach Menschen lernen müssen, dass sie grundsätzlich ihren Augen und ihren Ohren misstrauen. Dritter Punkt, der mir, glaube ich, die größten Sorgen macht. Und das ist der Punkt der krassen Machtkonzentration, die KI mit sich bringt.

Das haben wir ja schon in den letzten 10, 15 Jahren bei generell Tech und Social Media erlebt, dass einfach eine sehr mächtige Technologie in der Hand von sehr wenigen Konzernen und damit auch Menschen an der Spitze dieser Konzerne ist. Und das wird durch KI noch extremer. Das hier ist jetzt eben Sam Altman. Der hat, das finde ich, ein sehr vielsagendes Zitat einmal verwendet. Er hat jemand anderen zitiert, es aber auf seinem Blog veröffentlicht und sich dem angeschlossen: „Erfolgreiche Menschen gründen Firmen, erfolgreichere Menschen gründen Länder, die erfolgreichsten Menschen gründen Religion.“ Und genau das ist OpenAI. OpenAI ist am Ende eine Religion und er ist so der Sektenführer und das macht mir Sorgen.

Womit ich noch weniger anfangen kann und wen ich für noch problematischer halte, das ist dieser Herr hier, eben Mark Zuckerberg, der Chef von Meta, das ja auch sehr viel mit KI gerade macht, der sich in dem vergangenen Jahr einfach völlig schamlos Donald Trump unterworfen hat, den ich für einen ziemlich gewissenlosen Opportunisten halte, bei dem es mir auch kein gutes Gefühl macht, wenn er jemand ist, der so mächtige Technologie kontrolliert. Und genau bei diesem Herrn hier muss ich, glaube ich, gar nicht so viel sagen. Das nimmt er mir mit seinen öffentlichen Auftritten ja schon ab. Aber die Tatsache, dass Elon Musk mit xAI, einem weiteren KI-Unternehmen, auch einen sehr mächtigen, mit Dutzenden Milliarden Dollar bewerteten KI-Konzern gerade gegründet und groß gemacht hat. Das macht mir deutlich größere Sorgen als möglicherweise intelligente und wild werdende Maschinen.

Und um jetzt nicht den Vortrag mit Elon Musk zu enden, möchte ich mit zwei, finde ich, deutlich angenehmeren und vernünftigeren Menschen aufhören. Das sind zwei Wissenschaftler, unter anderem aus Princeton, also sehr renommierte Forscher in den USA, die ein Buch geschrieben haben und immer wieder über „AI as normal technology“ sprechen. Und das finde ich die hilfreichste Perspektive auf KI, dass man sagt, es ist keine magische Technologie, es ist vielleicht eine mächtige Technologie, aber es ist am Ende einfach eine Technologie, wie viele andere auch,

die man kontrollieren kann und regulieren kann. Wir sind KI nicht ausgeliefert. Es ist keine so unaufhaltsame Revolution, die einfach von sich aus über die Menschheit kommt, sondern wir können darauf hinwirken, KI so zu gestalten, dass sie Menschen dient und Menschen nicht darunter leiden müssen.

24:10 – Fragen und Abschluss

Simon Berlin

Und damit bin ich jetzt erst mal durch und genau, ich schaue mir jetzt die Fragen an und beende die Präsentation, damit ich im Vollbild zu sehen bin. Sekunde. Ich müsste genau ... Ah, hier. Jetzt bin ich wieder zu sehen. Oh, okay.

[Lachen]

Jetzt hat mir gerade meine Kollegin was rübergeschoben und ich sehe, dass im Chat bislang keine Fragen aufgetaucht sind. Das heißt, ich kann jetzt nur noch mal dazu appellieren. Jetzt wäre noch Zeit, Fragen zu stellen. Ich warte noch so, ich weiß nicht, eine Minute, falls noch jemand irgendwas wissen möchte, entweder zu den Sachen, die ich angesprochen habe oder generell zum Thema KI. Dann gebe ich mir Mühe, das zu beantworten. Einfach in den Chat schreiben. Und falls dem nicht so ist, dann hoffe ich, dass ich alle möglichen Fragen beantwortet habe, dass Sie ein bisschen was gelernt haben und wünsche denjenigen, die sich jetzt verabschieden, schon mal einen schönen Abend.

[Stille]

Ah

[lacht]

– Nur sehr entfernt zu KI. Der Newsletter, den ich schreibe, heißt Social Media Watchblog. Ich schreibe den zusammen mit einem Freund und Kollegen, Martin Ferenzen. Erscheint zweimal pro Woche und wir machen das seit 13 Jahren und versuchen so, früher eben, merkt man auch im Namen, haben wir vor allem über Social Media geschrieben, mittlerweile viel über KI. Es geht so ein bisschen um die Zusammenhänge zwischen Politik, Gesellschaft und Technologie. Und was bedeutet das eigentlich, dass es früher diese großen Plattformen gab und jetzt KI gab? Nachdem sonst bislang keine weiteren Fragen aufgetaucht sind, sage ich jetzt noch mal danke für die Aufmerksamkeit und würde mich jetzt auch verabschieden und wünsche allerseits einen schönen Abend.