

GROSSE SPRACHMODELLE

Machine Learning als Lese- und Schreibermöglichkeit

von HANNES BAJOHR

Große Sprachmodelle in der Rubrik «Werkzeuge» zu besprechen, heißt bereits, Stellung zu beziehen. Denn die Diskussion darum, ob *large language models* (LLMs) *tools* oder *agents* sind, ist keineswegs ausgefochten.¹ Letztere Position muss gar nicht einmal so gemeint sein wie die Rede von der *artificial general intelligence*, die OpenAI-Chef Sam Altman als Endziel der technischen Entwicklung stets auf den Lippen führt. Es könnte ja schon genügen, LLMs als Partner in einer «künstlichen Kommunikation» zu betrachten, die hinreichend unvorhersehbar sind, um die «doppelte Kontingenz» kommunikativen Verhaltens herzustellen.² Und auch die Schreibszenenforchung sollte sich in der Verabschiedung von der Idee, Texte seien notwendig von Menschen hergestellt, mit der umgekehrten Annahme, Maschinen könnten ein «Subjekt des Schreibens» sein oder daran Anteil haben, durchaus anfreunden können.³

Dass ich große Sprachmodelle dennoch als Werkzeuge behandle, liegt aber auch daran, dass sie im Moment vor allem noch als solche verwendet werden. Ich will hier zwei Aspekte dieser Werkzeughaftigkeit herausgreifen. Der erste zeigt sich bereits vor der Herstellung jenes Textes, der für gewöhnlich als das eigentliche Produkt der Schreibarbeit gilt. Bei Übersetzungsdiensten wie DeepL oder Sprachtranskriptoren wie Otter.ai ist das evident. Eine viel seltener bemerkte Funktion scheint mir die *Ermöglichung des Lesens* zu sein, nämlich durch die Zusammenfassung von bereits Geschriebenem. Bitte ich etwa Claude, das LLM der Firma Anthropic, einen Essay auf seine Grundargumente zu reduzieren, kann ich eine Idee davon bekommen, ob es sich lohnt, ihn überhaupt zu lesen. Eine solche «Synoptierbarkeit» scheint eng an Textgenres und Disziplinengrenzen gebunden zu sein. Artikel der empirischen Sozialwissenschaft, informatische Whitepapers, sogar überlange Lexikoneinträge lassen sich oft hinreichend gut zusammenfassen. Dagegen sind etwa ein Lacan-Seminar oder auch auf still Mitgemeintes angelegte Alltagskommunikation und die meisten

¹ Leah Henrickson: Tool vs.

Agent: Attributing Agency to Natural Language Generation Systems, in: Digital Creativity, Bd. 29, Nr. 2–3, 2018, 182–190, doi.org/10.1080/14626268.2018.1482924.

² Elena Esposito, Kommunikation mit unverständlichen Maschinen, Wien, Salzburg 2024, Kap. 5 und 7.

³ Moritz Hiller: Es gibt keine Sprachmodelle, in: Davide Giuriato, Claas Morgenroth, Sandro Zanetti (Hg.): Noten zum «Schreiben», Paderborn 2023, 279–285, hier 280; vgl. auch den von Moritz Hiller und mir herausgegebenen Text+Kritik-Sonderband Das Subjekt des Schreibens, München 2024 [im Erscheinen].

Arten von Poesie selten so reduzierbar, dass das nicht rein Propositionale in ihnen erhalten bleibt. Wer dunkel schreibt, so könnte man sagen, wird in Zukunft vielleicht immer noch nicht gelesen, aber nicht nur nicht von Menschen, sondern auch nicht von Maschinen. Das ist LLMs aber kaum vorzuwerfen. Wie alle Werkzeuge können sie nicht auf jede Domäne angewendet werden, auch wenn mit ihnen als Hammer alles wie ein Nagel aussehen will.

In ihrer Lesefunktion sind LLMs also zunächst einerseits simples Hilfsmittel. Wollen die Digital Humanities dem *great unread* auf Distanz Herr werden, können sich die restlichen Geisteswissenschaften dem Versprechen hingeben, *machine learning* Sorge nun auch auf dem Nahfeld des Lesens für Übersicht – oder gestalte, einem Stoßseufzer Arno Schmidts eingedenk, die Divergenz von Lebenszeit und Lesezeit doch ein bisschen weniger steil.⁴ Andererseits bietet KI die Lösung für ein Problem, für das sie selbst verantwortlich ist: Wenn große Sprachmodelle, wie Matthew Kirschenbaum meint, eine «Textpokalypse» auslösen und das Web und unser Leben mit synthetischem Text überschwemmen, sind sie über kurz oder lang auch das Mittel, dieser Flut wieder Herr zu werden, eben durch Zusammenfassungen.⁵ Wie realistisch Kirschenbaums Vorhersage ist, zeigte erst kürzlich eine Untersuchung, der zufolge die Wendung *to delve into* (in ein Thema eintauchen), die zu den Lieblingsphrasen von ChatGPT gehört, in Artikeln auf PubMed heute 10- bis 100-mal häufiger auftaucht als noch vor zehn Jahren.⁶

Das ist in zweierlei Hinsicht bemerkenswert. Erstens wird die Zusammenfassung als Leseermöglichung hier wieder zum Schreibverfahren, nämlich dann, wenn sich Autor*innen ihre Abstracts generieren lassen. Und zweitens kehrt sich, lasse ich mir per LLM einen Text resümieren, für dessen Herstellung bereits LLMs verwendet wurden, das in der Signaltechnik standardmäßige Verhältnis von Kompression und Dekompression um. Denn im Normalfall geht es bei Informationsübertragung darum, die Redundanz einer Botschaft angesichts potenzieller Rauschquellen so hoch wie nötig, angesichts begrenzter Kanalkapazitäten zugleich aber so gering wie möglich zu halten. Bei hochredundanter natürlicher Sprache ist für die Übertragung über einen Kanal eine Kompression möglich, auf die dann auf der Empfänger*innenseite die Dekompression folgt (>–<). In der genannten Bewältigung der «Textpokalypse» träte aber gerade der umgekehrte Fall ein: Die Nachricht wäre auf Sender*innen- wie Empfänger*innenseite bereits komprimiert, die Dekompression würde zum Übertragungscodec des Kanals (⇔).

In diesem Fall sind LLMs mehr als nur simples Schreibhilfsmittel, sondern, wie man mit Arnold Gehlen, Walter Benjamin und Hans Blumenberg sagen könnte, Werkzeuge zur psychosensorischen Entlastung vom Textlich-Absoluten. Dass diese Beobachtung nicht ganz fantastisch ist, lässt sich an mehreren Stellen beobachten. So gibt es bereits Funktionen in kommerziellen Programmen (etwa in der aktuellen Office Suite), die aus einer Reihe von Stichworten ganze E-Mails formulieren, während sie empfangene Nachrichten in Stichworten

⁴ Arno Schmidt: Julianische Tage, in: ders.: *Bargfelder Ausgabe*, Werkgruppe 3: Essays und Biographisches, Bd. 4: Essays und Aufsätze II, Zürich 1995, 87–92.

⁵ Matthew G. Kirschenbaum: Prepare for the Textpocalypse, in: *The Atlantic*, 8.3.2023, www.theatlantic.com/technology/archive/2023/03/ai-chatgpt-writing-language-models/673318 (8.6.2024).

⁶ Jeremy Nguyen: Are medical studies being written with ChatGPT?, Twitter/X, 30.4.2024, twitter.com/JeremyNguyenPhD/status/1774021645709295840 (8.6.2024); vgl. auch Alex Hern: How Cheap, Outsourced Labour in Africa is Shaping AI English, in: *The Guardian*, 16.4.2024, www.theguardian.com/technology/2024/apr/16/techscape-ai-gadgets-humane-ai-pin-chatgpt (8.6.2024). Eine frühere Studie kommt zu einem ähnlichen Ergebnis, vgl. Weixin Liang u. a.: Monitoring AI-Modified Content at Scale: A Case Study on the Impact of ChatGPT on AI Conference Peer Reviews, in: *arXiv*, 15.7.2024, doi.org/10.48550/arXiv.2403.07183.

7 James Betker u. a.: Improving Image Generation with Better Captions [Research Paper zu DALL-E 3], OpenAI, 19.10.2023, cdn.openai.com/papers/dall-e-3.pdf (8.6.2024).

8 Andrej Karpathy: The Hottest New Programming Language is English, Twitter/X, 24.1.2023, twitter.com/karpathy/status/1617979122625712128 (8.6.2024).

9 Präsidium der DFG: Stellungnahme des Präsidiums der Deutschen Forschungsgemeinschaft (DFG) zum Einfluss generativer Modelle für die Text- und Bilderstellung auf die Wissenschaften und das Förderhandeln der DFG, DFG, 21.9.2023, www.dfg.de/resource/blob/289674/ff57cf46c5ca109cb18533b21fba49bd/230921-stellungnahme-praesidium-ki-ai-data.pdf (8.6.2024).

10 Jenifer Becker: Dear GPT-3. Collaborative Writing with Neural Networks, in: Eckart Voigts u. a. (Hg.): Artificial Intelligence – Intelligent Art? Human-Machine Interaction and Creative Practice, Bielefeld 2024, 189–202, doi.org/10.14361/9783839469224-013.

11 Juan S. Guse: Das kombinatorische Seekuh-Prinzip, in: Hannes Bajohr, Ann Cotten (Hg.): Schreiben nach KI, Berlin 2024 [im Erscheinen].

12 Vgl. Simon Roloff, Hannes Bajohr: Digitale Literatur zur Einführung, Hamburg 2024 [im Erscheinen], und Hannes Bajohr: Writing at a Distance: Notes on Authorship and Artificial Intelligence, in: German Studies Review, Bd. 47, Nr. 2, 2024, 315–337.

13 Hannes Bajohr: (Berlin, Miami), Berlin 2023; Genauer zur Machart findet sich darin im Nachwort, ebd., 239–273.

zusammenfassen; Elaboration ist Schnittstelle zwischen Maschinen, nicht zwischen Menschen, die nur die reduzierte Version eines Textes erhalten. Ähnlich werden auch bei der Text-zu-Bild-KI DALL-E die Prompts der User*innen nicht mehr direkt dem Bildgenerator übergeben, sondern vom System erst ausgeschmückt und mit mehr Details versehen. Dieser *ornatus* hat, so die Beobachtung der OpenAI-Ingenieur*innen, bessere Resultate zur Folge, bleibt für User*innen aber weitgehend unsichtbar.⁷ Damit ließe sich das Aperçu des Programmierers Andrej Karpathy, das auf die neue Macht natürlicher Sprache abhebt – «the hottest new programming language is English»⁸ –, reformulieren: «the hottest new transmission protocol is verbose English». Rhetorik ist nicht nur Code, sondern auch Codec.

Die Ermöglichung des Lesens spielt, das sollte deutlich geworden sein, zweitens bereits in die *Ermöglichung des Schreibens* hinein. KI-generierte Abstracts sind nur ein Beispiel für jene Genres, die man Texte ohne *jouissance* nennen könnte. LLMs versprechen, die falsche Schreibzeit für lebensqualitätsvermiesende Routinen wie Antragsprosa, Verwaltungskommunikation und Abschlussberichte an die Maschine auszulagern, um der Forschung mehr echte Schreibzeit einzuräumen. Dass die DFG derartiges inzwischen ausdrücklich erlaubt, solange der Einsatz von generativer KI ausgewiesen wird, ist sicher nicht zuletzt der Einsicht geschuldet, dass immer mehr Forscher*innenleben auf ein *œuvre cachée* abgelehnter Anträge vergeudet werden, die kein akademisches Publikum mehr zu Gesicht bekommt.⁹

Die Unterscheidung zwischen falscher und echter Schreibzeit geht aber womöglich da fehl, wo auch das «echte» Schreiben ohne Lustgewinn betrieben wird. Deshalb mag es interessanter sein, sich hier der Literatur als Verdachtsfall von *jouissance*-erfüllter Textproduktion zuzuwenden. Die Schriftstellerin Jenifer Becker hat ihren Arbeitsmodus mit GPT-3 zunächst als kollektive Ideenfindung umrissen, dem *writers' room* von Fernsehserien gleich, in dem Vorschläge hin und her gespielt, aufgenommen und wieder verworfen werden können.¹⁰ Wie er selbst zu einer ähnlichen Einschätzung gekommen ist, beschreibt der Autor Juan S. Guse in zwei Schritten: Habe er ChatGPT zunächst für eine «Zusammenarbeit *ex negativo*» verwendet, um «stochastisch ausgetretene Pfade» zu vermeiden – kommt die KI auf denselben Gedanken wie ich, ist der Gedanke schlecht –, so sei seine Praxis mittlerweile eher «mäeutisch», wenn er die Vorschläge des Programms nun auch positiv aufnehme.¹¹

Noch scheint es aber um Ideen-, nicht um Textgenerierung zu gehen. Letzteres ist im Augenblick am ehesten bei digitaler Literatur der Fall, die besonders offen ist für Schreibverfahren, die in einer gewissen Distanz zu ihrem*ihrer Urheber*in stehen, und zugleich keine konventionellen Gattungskonventionen bedienen muss.¹² In dieses Genre gehört auch mein Roman (*Berlin, Miami*), der mit von mir auf Gegenwartsliteratur feinabgestimmten offenen LLMs namens GPT-J und GPT-NeoX entstanden ist.¹³ Der Text war dabei nicht Resultat eines Prompts («Schreibe mir einen Roman!»), sondern

kam über wiederholte Satzvervollständigung zustande. Um herauszufinden, ob und auf welche Weise ein solches Modell zur Narration fähig ist, griff ich dabei so wenig wie möglich ein und ließ es oft einfach laufen. Die Erfahrung des Schreibens war daher zwar eine besonders distanzierte. Dennoch erfuhr ich das LLM in diesem produktiven Moment keineswegs als Agenten. Eher hatte ich den Eindruck, es sei der Text, der etwas wollte und in eine bestimmte Richtung drängte. Man surft auf Ideen, die aber regelmäßig im Absurden oder Abschweifenden versanden. Aus diesem Grunde griff ich an manchen Stellen aktiv ein, ließ verlorene Fährten erneut aufnehmen oder führte Themen ein – oft nur mit einem Wort –, über die ich mehr wissen wollte. So traf ich etwa auf den vom Modell erfundenen und nur vage beschriebenen «Kieferling», der zusammen mit dem ebenso unklar bleibenden «Teichenkopf» sein Unwesen trieb. Oder ich lernte etwas über den Prozess der «Diagonalisierung», dem die von akuter urbaner Verlotterung betroffene Stadt Miami ausgesetzt war, während sie unter der Ägide der Äää-Agentur zu leiden hatte. Das war alles so interessant, dass ich hin und wieder nachhakte, wenn diese Themen zu verschwinden drohten, aber immer offen für das blieb, was da noch kommen mochte. Die Schreibtätigkeit bestand dann vor allem im Wechsel zwischen surfendem Laufenlassen und kleinen *nudges* – dem, was beim Reiten mit dem schönen Wort «Schenkelhilfe» bezeichnet wird, von der es sowohl die «vorwärtstreibende» wie auch die «verwährende» Variante gibt.¹⁴ Zwischen diesen Metaphern, Surfen und Reiten, bewegt sich – jedenfalls bei mir, jedenfalls im Moment noch – die Erfahrung des Schreibens mit KI. An dieser Stelle verwischt sich auch die Unterscheidung zwischen *tool* und *agent*. Wie man ein Surfbrett, mag man es noch so sehr als Auswuchs des eigenen Körpers erfahren, noch nicht als Agenten bezeichnen will, so wäre es falsch, von einem Pferd als einem Werkzeug zu sprechen.¹⁵ LLMs sind am Ende womöglich ein Drittes, für das man erst eine Praxis und einen Namen finden müsste.

¹⁴ Deutsche Reiterliche Vereinigung (Hg.): *Richtlinien für Reiten und Fahren*, Bd. 1: Grundausbildung für Reiter und Pferd, Warendorf 2014, 83. Es gibt zudem noch die «vorwärtsseitwärtstreibende» Schenkelhilfe, die in diesem Kontext vielleicht ein besonders ergebnisoffenes Eingreifen meinen könnte.

¹⁵ Ich kann weder reiten noch surfen, aber so stelle ich mir das vor.