

Protokoll 10

Paul Heinicker

Post-statistisches Assoziieren

Wie orten und werten wir die Antworten von ChatGPT? Für diese Frage scheint im öffentlichen Diskurs um generative KI allgemein oder die öffentliche Verfügbarmachung von *Large Language Models*, wie etwa durch OpenAI, wenig Platz. Denn ihre Auswirkungen und Potenziale werden vornehmlich an den Rändern eines Spektrums diskutiert. Einerseits hätten wir es schlicht mit einer mit Milliarden von Daten aufgebauchten statistischen Autovervollständigung zu tun, deren zu häufig falsche Ergebnisse eine produktive Anwendung ausschließen und allenfalls zur Belustigung in sozialen Medien geteilt werden. Andererseits seien die makellosen Ergebnisse dieser generativen Modelle der Beweis für das automatisierte Ende einiger Berufszweige, wenn nicht gleich der öffnende Pfad zur *Artificial General Intelligence*.

Irritiert von dieser klaren Entscheidung, suche ich im Schatten dieser zwei Säulen nach Möglichkeitsräumen für meine eigene Anwendungspraxis als Datengestalter. Und die braucht Hilfe denn gerade programmiere ich wieder. Normalerweise bedeutet das ein Hangeln durch Suchanfragen wie "how to merge multiple arrays in javascript" auf Stack Overflow, um aus meinem eigenen Flickenteppich, durchzogen von Erinnerungslücken, mit der Hilfe von fremdem Code ein halbwegs stabiles Gerüst zu bauen – business as (coder) usual. Gerade weil ChatGPT auf einem Datensatz namens *Common Crawl* kalibriert wurde, der damit gerühmt wird, einen großen Teil der öffentlich zugänglichen Webseiten zu beinhalten, finden sich auch viele Inhalte aus solchen Austauschplattformen für Programmierer:innen, wie eben Stack Overflow. So findet ChatGPT tatsächlich auf solch formalisierte Fragen ausführliche Antworten mit Code-Beispielen und Erklärungshilfen.

Auf meine obige Frage, wie mit der Kombination meiner Datenstrukturen umzugehen ist, hat ChatGPT direkt eine anwendbare Antwort getreu meiner abgefragten Variablen: eine *foreach*-Schleife. Auf die Idee kam ich aber schon selbst und schien mir viel zu umständlich für mein wesentlich simpleres Problem. Allgemein

funktioniert der generierte Code von ChatGPT für mich nur in sehr konkreten Fällen und losgelöst von meiner Programmierlogik. Die Integration der vermeintlichen Lösungen braucht dann immer noch ein geschultes Auge und zumindest so viel Expertise, den generierten Code als nicht funktional einordnen zu können. Hier muss die Zusammenarbeit mit ChatGPT aber nicht aufhören. Salopp frage ich also direkt hinterher "is there another way do to this?". There is: im zweiten Anlauf bekomme ich meine vergessene *flatMap*-Methode, die es mir ohne Schleifen erlaubt meine Datenstrukturen direkt miteinander zu kombinieren.

Ab diesem Moment ändert sich meine Erwartungshaltung. Es geht von nun an weniger darum, Prompts zu stellen, die möglichst genau mein Problem beschreiben und ich mir dahingehend eine Effizienz in der Lösung dieser erhoffe. Vielmehr versuche ich einen ergebnisoffenen Raum zu gestalten, in dem sich andere Formen der Fragestellung ergeben: "how would this code look like if...?" oder "can you explain this code in more depth?". Viel häufiger benutze ich ChatGPT ohne Erwartung einen anwendbaren Code zu bekommen, sondern lasse mir vielmehr unterschiedliche Angebote machen bzw. die Idee dahinter erklären. So kommt es oftmals vor, dass ich gerade durch die vielen nicht funktionierenden oder nicht integrierbaren Vorschläge auf Ideen für ein anderes Ergebnis komme. Der Mehrwert von ChatGPT liegt dann eher im Freilegen von Assoziationsketten durch oder vielmehr trotz der Resultate der automatisierten Wahrscheinlichkeitsverteilung.

Dieses Denkmodell aus meiner Programmierpraxis mit ChatGPT versuche ich auf meine datengestalterische Arbeit insgesamt übertragen. Gerade vor der Datenprozessierung und der konkreten Umsetzung einer Datenordnung in eine Visualisierung stehen viele konzeptionelle Fragen, die alle einer Entscheidung bedürfen. Wie begegne ich einem Visualisierungsgegenstand? Für welche Metaphern und Symbole entscheide ich mich? Welche Optionen schließe ich aus? Bei solchen Ideenfindungen geht es auch wieder um Assoziationsketten. Gestalter:innen nähern sich über Kreativitätstechniken langsam ihrem Entwurfsideal. Es geht darum viele Ideen zu generieren und ebenso viele wieder zu verwerfen, bis sich eine Iteration als zweckmäßig erweist. Wie kann sich ChatGPT in einem solchen Entwurfsprozess positionieren?

„Give me 20 ideas for a visualisation of a time-based dataset“. In wenigen Sekunden habe ich eine Fülle von formalen Ideen für einen spezifischen Datensatz. Bei genauerer Betrachtung arbeiten sich diese aber an den üblichen Konventionen und Standards bezüglich Lesbarkeit und Effizienz ab, sprich Varianten von Liniengraphen, Flächendiagrammen oder Kalenderdarstellungen. Gemäß seiner Datengrundlage bekomme ich bei ChatGPT vor allem die Ideen, die sich als erfolgreich im Web etabliert haben und vorrangig geteilt werden. Durch das statistische

Assoziieren bekomme ich ein Ranking der Verbreitung etablierter Visualisierungsideen – Konsens statt Innovation.

Erinnert hat mich das an eine Übung, die im Designstudium oft Anwendung findet. Dabei werden bei einer Ideenfindung die ersten Einfälle niedergeschrieben und dann als zu naheliegend und einfach für den Entwurfsprozess gebrandmarkt. Man lernt mit vielen Ideen zu haushalten und seiner eigenen Intuition zu misstrauen. Und dafür eignet sich die Integration von ChatGPT bestens. Ich kann wie dargestellt die betretensten Pfade eines Gestaltungsproblems ausfindig machen und mich durch die Methode des Ausschlusses meinem Entwurfsideal nähern. Es birgt ein großes Potenzial mehrere Iterationsstufen im Entwurf kürzen zu können. Erneut sind es wieder nicht die statistischen Vorhersagen, die mein Problem faktisch lösen, aber gerade durch ein informiertes Gegenlesen eröffnen sie wiederum einen tiefergehenden Gestaltungsprozess.

—

Beide Beispiele meiner Praxis mit ChatGPT in der Programmierung und der Entwurfssynthese eint eine Abkehr von der Vorstellung, deren Ergebnisse ans Ende einer Suche nach Erkenntnissen stellen zu wollen. Was wenn die Zweckmäßigkeit generativer Modelle gar nicht so sehr darin liegt gute oder wahre Resultate hervorzubringen, sondern ihr Potential genau in ihrer statistischen Fragilität liegt? An welchen Orten können wir generative Modelle denken, wenn wir sie von den Verantwortungen lösen, die sie strukturell gar nicht einlösen können? Gerade abseits projizierter Szenarien drohender Vollautomationen sehe ich gerade in der Geistes- und Entwurfsarbeit Potenzial in der angedeuteten Assoziationshilfe. Hier besteht die Chance für ein kritisches Handlungswissen mit und über statistische Modelle hinaus. Meine skizzierte ChatGPT-Praxis fasse ich dahingehend als ein „post-statistisches Assoziieren“.