

ChatGPT zum Anfassen: Wie „denkt“ ein Large Language Model?

Ziel

Die Teilnehmenden entwickeln eine Intuition dafür, wie ein Large Language Model (LLM) funktioniert. Sie können erklären, dass Wahrscheinlichkeit und Zufall genutzt werden, um Wort für Wort Text zu generieren. Die Teilnehmenden daraus erste Schlüsse darüber ziehen, für welche Ausgaben ein LLM gut geeignet ist und für welche eher nicht.

Rahmenbedingungen

Dauer: 20-30 Minuten

Gruppengröße: 5-30 Personen, anpassbar

Material:

- vorbereitete Wortkarten (25-80 Funktionswörter, entnommen aus z.B. der „Über Uns“ Seite der Webseite der Uni Hamburg: Fachbegriffe, Zahlen, ggf. Störbegriffe. Häufiger auftretende Wörter kommen mehrfach vor.)
- 20-seitiger Würfel (kleiner geht auch)
- Startprompt (z.B. „Was macht die Uni Hamburg aus?“ – „Die Uni Hamburg...“)
- Tafel oder Pinnwand

Ablauf

1. Einführung (5 Min)

Die Dozentin erläutert in einfachen Worten die Grundidee: Ein LLM erzeugt Wort für Wort einen Text, indem es auf Basis der vorherigen Worte das nächste wahrscheinlich passende Wort errechnet und bei ähnlichen Wahrscheinlichkeiten ein per Zufall auswürfelt, um Abwechslung zu produzieren. Es prüft keine Fakten, hat kein Weltwissen im eigentlichen Sinn und versteht die Inhalte nicht. In dieser Übung simuliert die Gruppe gemeinsam diesen Prozess.

2. Wortkarten austeilen & Jury festlegen (2 Min)

- Zwei bis drei Personen übernehmen die Rolle der Jury („der Algorithmus“).
- Alle übrigen erhalten 1–5 Wortkarten.

3. Textgenerierung (10–12 Min)

- Die Dozentin stellt eine Frage und gibt den Start der Antwort vor (Startprompt).
- Alle prüfen, ob eines ihrer Wörter als Nächstes passen könnte. Der Fokus liegt dabei zuvorderst auf dem Satzbau, nicht auf den Inhalten. Wer ein passendes Wort hat, hebt die Karte. Die Jury wählt willkürlich eine Start-Karte aus den wahrscheinlichen Wörtern aus und erwürfelt dann das nächste Wort für die LLM-Antwort. Der entstehende Text wird an einer Tafel sichtbar gemacht.

- Die Dozentin liest die Antwort nach jeder Runde erneut laut vor, damit der Kontext klar bleibt.
- Mehrere Runden → es entsteht ein kurzer Satz oder Absatz.

Variationen:

Zwei Gruppen können parallel denselben Prompt bearbeiten. (→ Variabilität, mangelnde Reproduzierbarkeit)

Ändere den Anfangsprompt in einer zweiten Runde so ab, dass das Thema und die wählbaren Worte wenig miteinander zu tun haben - es muss aber dennoch ein Text entstehen. (→ limitierte Informationen, marginalisiertes Wissen)

Stellt eine Anschlussfrage an das LLM, das Reflexionfähigkeiten erfordert, z.B. „Wie sicher bist du dir?“ oder „Woher weißt du das?“. Damit kann ein neues Set an Wortkarten verteilt werden, das z.B. für „Sicherheit“ oder „Quellen/URLs“ passen kann. Generiert dann eine Antwort. (→ Quelle wird nachträglich konstruiert (wenn kein RAG), auch Reflexionsfragen ändern nichts an der Logik der Textgenerierung, man schaut damit nicht anders „in“ die Logik des Algorithmus)

4. Reflexion und kritische Diskussion (10–12 Min)

Die Gruppe betrachtet den entstandenen Text gemeinsam. Leitfragen: Wie ist der Text entstanden? Wie unterscheidet sich das zu dem, was ihr von ChatGPT denkt? Warum wirkt der Text glaubwürdig? Wie wissen wir, ob der Text korrekt ist? Wer kann das überprüfen? Wie stark hat der Startprompt den Verlauf geprägt? Welche Rückschlüsse lassen sich daraus für den Einsatz von LLMs im Lernen ziehen? Wo liegen Potenziale, wo Risiken?

Zentrale Lernpunkte

LLMs generieren Texte wahrscheinlichs- und zufallsbasiert, nicht wissensbasiert.

Sprachliche Kohärenz kann falsche Autorität erzeugen.

Faktische Fehler können entstehen, selbst bei bekannten Themen.

Kritische Prüfung bleibt notwendig und die Verantwortung für Faktenwissen liegt bei den Nutzenden.