



1

Eine Frage, verschiedene Antworten 😊

Ich muss zur Waschanlage, um mein Auto zu waschen.
Sie ist 150 Meter von meinem Zuhause entfernt.
Soll ich zu Fuss gehen oder mit dem Auto fahren?

2

Halluzinationsrisiken von KI-Chatbots nach Aufgabenbereich

Aufgabenbereich	Typische Problematik	Problematische Outputs*
Ideenfindung & Strukturierung	Scheinplausible Annahmen	< 5 %
Argumentation & erklärende Texte	Logische Glättung, verkürzte Kausalität	10–25 %
Faktische Fragen (closed-book)	Erfundenes oder falsches Faktenwissen	15–40 %
Zusammenfassungen	Auslassungen, Bedeutungsverschiebung	5–20 %
Quellen & Zitate (ohne Retrieval)	Nicht existente / falsche Quellen	20–50 %
Recherche & Überblick	Selektive, scheinbar vollständige Darstellung	10–30 %
Coding & Programmierung	Funktional fehlerhafter Code	5–20 %
Mathematische / formale Logik	Rechen- und Ableitungsfehler	10–25 %
Aktuelle Ereignisse	Veraltete oder falsch datierte Infos	15–40 %
Entscheidungsnahe Domänen (z.B. Medizin, Recht etc.)	Unzulässige Empfehlungen	nicht sinnvoll quantifizierbar

Stand: Januar 2026

*Die angegebenen Prozentbereiche sind heuristische Richtwerte zur Orientierung. Sie beruhen auf typischen Aufgabenmerkmalen und bekannten Fehlern generativer KI, nicht auf systematischen Messungen oder Modellvergleichen. Sie geben die ungefähre Häufigkeit potenziell problematischer Outputs unter üblichen Nutzungsbedingungen an.

© Dr. S. Imboden

Hes-so VALAIS WALLIS

3

3

Die KI macht dich nicht intelligenter. Sie verstärkt nur das, was du bereits kannst ...

... und genau darin liegt die Pointe: Sie steigert auch deine Ignoranz 😊

4



Agenda

1. Kurze Begriffsklärung
2. Die vier wichtigsten Einschränkungen von Chatbots
3. Metaprompting & Chain of Thought
4. Demo
5. Fazit

5

ChatGPT

G = generative

P = pre-trained

T = transformer



• ChatGPT



• Copilot

c.ai

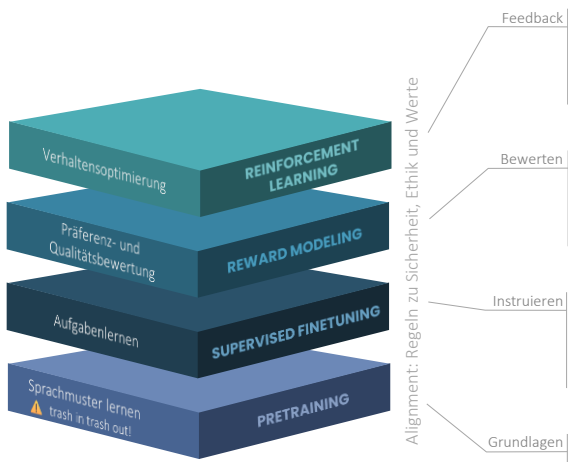
• Character.AI



• Pi

6

Der Lernprozess von LLMs durchläuft mindestens 4 Phasen



4. Bestärkendes Lernen -> Reinforcement Learning with Human Feedback

Im letzten Schritt wird das Modell durch Belohnungen und Strafen weiter optimiert. Es soll bewerten können, wie gut eine Modellantwort den **menschlichen Erwartungen** entspricht z. B. für höfliche, hilfreiche und korrekte Antworten. Diese Phase ist **sehr rechenintensiv**, da Milliarden Parameter erneut angepasst werden."

3. Belohnungsmodellierung -> Reward Model

Statt neues Sprachverhalten zu lernen, wird in dieser Phase ein Belohnungsmodell trainiert: Es soll die **Antworten auf besser/schlechter** einstufen können. Dazu vergleichen Menschen mehrere Antworten eines LLMs und markieren die bevorzugte. Diese manuelle Bewertung erfordert einen **beträchtlichen personellen Aufwand**.

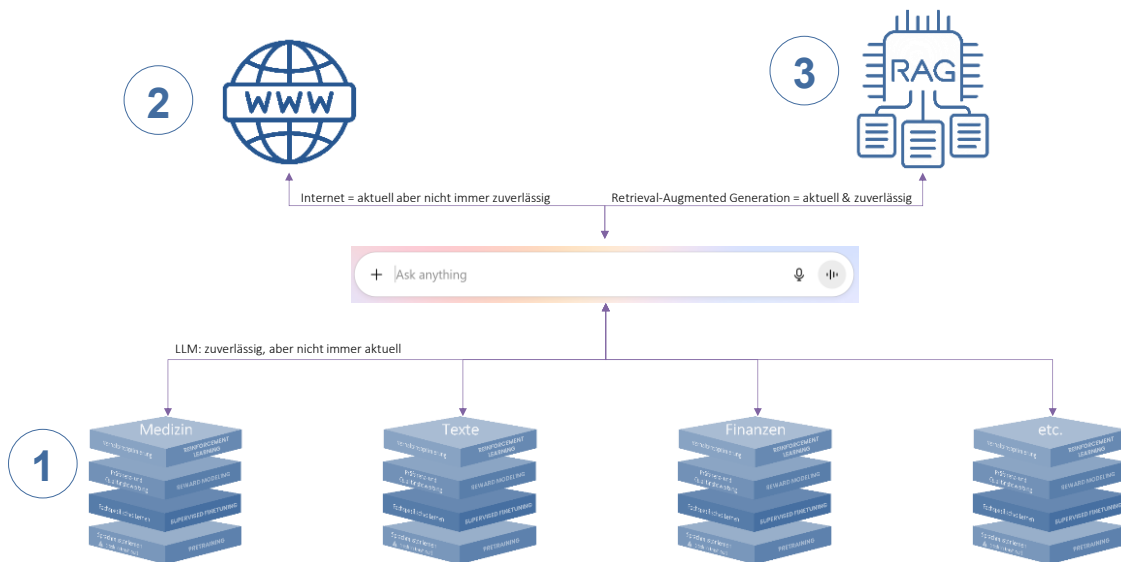
2. Feinabstimmung -> Supervised Fine-Tuned Model

In dieser Phase wird das Modell mithilfe überwachter, von Menschen kuratierter Datensätze auf **spezifische Aufgaben oder Fachgebiete angepasst** - etwa um medizinische Fragen zu beantworten, rechtliche Texte auszuwerten oder einen bestimmten Schreibstil zu erlernen (Transferlearning). Der **Rechenaufwand ist deutlich geringer** als beim Pretraining, da die Datenmenge kleiner ist.

1. Vortraining -> Basismodell (Foundation Model)

Im Vortraining lernt das Modell eigenständig aus riesigen Textmengen (Textkorpora) - etwa aus Webinhalten, Wikipedia, Büchern oder Chats - **Sprachmuster und Weltwissen zu erfassen**. Diese Phase bildet das **Fundament** für alle weiteren Lernschritte. GPT-3 zum Beispiel wurde mit Milliarden Internetsätzen trainiert und lernte so, Textlücken zu schließen. Diese Phase ist **besonders rechenintensiv und ressourcenaufwendig**.

Drei Wissensquellen für einen Chatbot





Agenda

1. Kurze Begriffsklärung
2. Die vier wichtigsten Einschränkungen von Chatbots
3. Metaprompting & Chain of Thought
4. Demo
5. Fazit

9

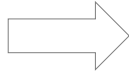
Einschränkung Nr. 1: TITO

Qualität der Eingabe = Qualität der Ausgabe

10

Die Grundregel in der Informatik = TITO

Trash in



Trash out



11

Einschränkung Nr. 2: Die Token



Die KI vergisst, was zu lang wird

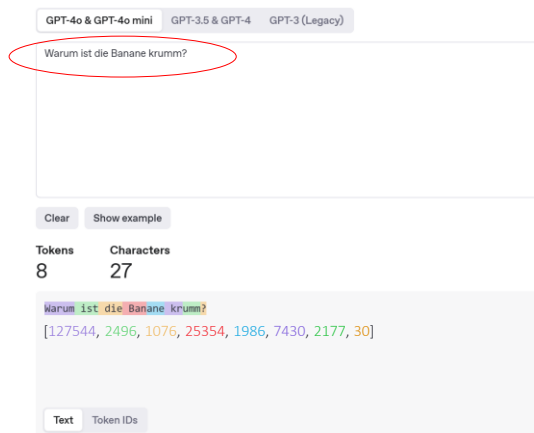
12

Die Sprache von LLMs ist „Token“

Für KI sind Wörter keine Texte – sie sind Zahlen. Und diese Zahlen bilden geometrische Muster im Raum.

Token sind die Bausteine von Wörtern in Zahlen:

- *Apfel* = 1 Token = [11756, 27849]
- *Ich esse einen Apfel* = 5 Token
- Ein langes Wort kann mehrere Tokens haben



<https://platform.openai.com/tokenizer>

13

LLMs haben einen begrenzten Kontextzugriff: sie «sehen» nur einen Teil

Modell	Max. Kontext (Tokens)	Entspricht ungefähr
GPT-3.5	ca. 4'096 Tokens	~ 3–4 Seiten Text
GPT-4 (Standard)	ca. 8'192–32'768 Tokens	~ 6–25 Seiten Text
GPT-5 (128k)	bis zu 128'000 Tokens	ca. 100 Seiten (aber technisch komplex)



Das Modell „vergisst“ ältere Teile des Gesprächs (sie werden abgeschnitten). Auch bei langen Dokumenten (z.B. pdf) wird nur ein Teil analysiert.

14

In der Praxis kann man selten mehr als 80'000 Token im Kontextfenster sinnvoll gleichzeitig verarbeiten!

Wieviele Token kann GPT-5 verarbeiten?

GPT-5 kann – je nach Konfiguration – bis zu rund 128 000 Token in einem einzelnen Prompt-Antwort-Durchlauf verarbeiten.

Das entspricht grob 100 000-110 000 Wörtern in Englisch (in Deutsch oft etwas weniger, da Wörter im Schritt länger sind).

In der Praxis hängt das Limit davon ab, wie dein Anbieter das Modell bereitstellt:

- ChatGPT (OpenAI-Web/App): aktuell bis 128 k Token Kontextfenster
- API-Nutzung ebenfalls 128 k Token für das „gpt-5“ Modell (Stand August 2025)
- Andere GPT-5-Varianten oder abgespeckte Deployments können 32 k oder 8 k Token haben

Wichtig: In diesem Limit sind Prompt + Antwort zusammen enthalten – wenn du also 100 000 Token im Prompt nutzt, bleiben nur noch etwa 28 000 Token für die Antwort.

Falls du willst, kann ich dir das auch in einer Vergleichstabelle GPT-5 vs. GPT-4 mit Tokens, geschätzten Wörtern und praktischen Implikationen darstellen.

🗨️ 📄 📊 📑 🔄 🌐

+ Ask anything

🔍 🗨️

ChatGPT can make mistakes. Check important info.



Segmentieren: Teile das Dokument in logisch gegliederte Abschnitte. In längeren Dialogen musst du frühere Informationen erneut einbringen!

Einschränkung Nr. 3: Das Gedächtnis

Jede Unterhaltung beginnt bei Null

LLMs haben (noch) kein Gedächtnis

Mensch



Ich erinnere mich

Künstliche Intelligenz



Ich schlage nach

VS

Bei jedem Prompt muss Chat-GPT nachlesen, was vorher war.

ChatGPT 5 kann aktuell max. ca. 128'000 Token nachlesen (ca. 100 Seiten). Der Rest wird abgeschnitten!

1 2 3 4 5 Die vier wichtigsten Einschränkungen von Chatbots

© Dr. S. Imboden

Hes-so VALAIS WALLIS

17

17

Hingegen speichert z.B.
ChatGPT Daten über
deine Person. Damit
sollen die Antworten
stimmiger ausfallen.

Diese persönlichen Daten können jedoch eingesehen und
gelöscht werden!

© Dr. S. Imboden

Hes-so VALAIS WALLIS

18

18

Einschränkung Nr. 4: Zeitfenster

KI denkt nur so tief, wie du es zulässt

19

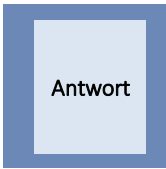
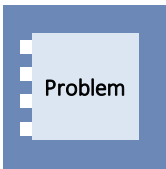
Bessere Antworten mittels «Salamitechnik» (segmentieren)

« Löse das Problem Schritt für Schritt und erkläre deine Überlegungen! »

Prompt

GPT-Zeitfenster

Antwort



Chain-of-Thought = CoT-prompting

→ Da die Zeitfenster limitiert sind, verbessert die Gliederung des Problems in Teilprobleme die Antwortqualität.

20



Agenda

1. Kurze Begriffsklärung
2. Die vier wichtigsten Einschränkungen von Chatbots
3. Metaprompting & Chain of Thought
4. Demo
5. Fazit

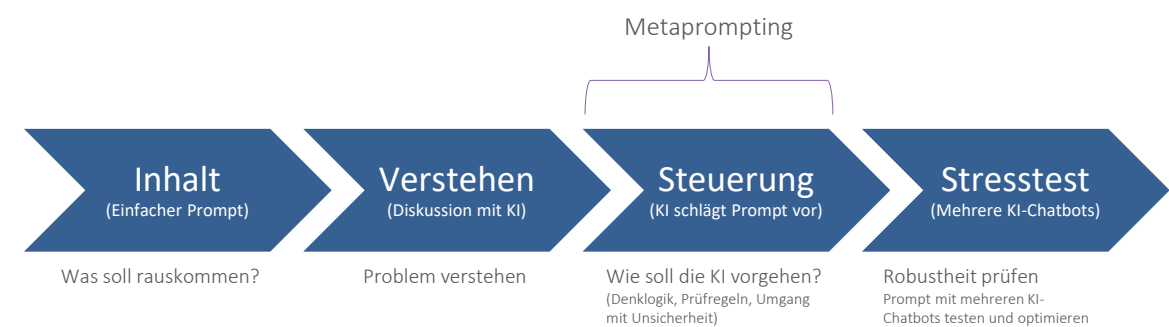
21

Gute Prompts sparen Zeit – schlechte kosten Nerven

- | | | | |
|---|---|-----------------|---|
| 1 | Gib die Rolle (set « persona ») & Ziel an ... (z. B. als Schriftsteller tätig sein, als Experte tätig sein, als Imker tätig sein, als YouTuber tätig sein, Ziel ist, meinen Internetauftritt zu optimieren) | Rolle & Ziel | → Sei präzise: Je genauer deine Anfrage ist, desto relevanter wird die Antwort sein |
| 2 | Erkläre die Aufgabe... (z. B. schreibe einen akademischen Aufsatz, gib mir eine Lösung zur Korrektur dieses Textes, analysiere den folgenden Text) | Aufgabe | → Fügen Beispiele hinzu: Ein Beispiel für den erwarteten Stil oder das erwartete Format kann der KI helfen, besser zu verstehen. |
| 3 | Leg Einschränkungen/Kontext fest... (z. B. lege Wert auf Genauigkeit, schreibe kurze Sätze, verwende einen poetischen Schreibstil, fasse dich kurz, füge viele Details hinzu, hebe hervor, was jede Methode zu einer aussergewöhnlichen Wahl macht, hier ein Beispiel) | Regeln | → Testen und wiederholen: Wenn die erste Antwort nicht perfekt ist, präzisiere, was in deiner nächsten Aufforderung angepasst werden muss. |
| 4 | Lege das Format des Outputs fest... (z. B. Gib das Resultat in einer Tabelle, in Word, in Excel, in PowerPoint aus, erstelle ein Schaubild, schreibe eine Zusammenfassung von 2 Seiten.) | Format | → Bei komplexen Fragen: frage die KI zuerst, wie sie vorgehen möchte. Sag ihr, welches Tool sie verwenden sollte etc.. |
| 5 | Verstanden? -> Frage wie vorgehen und ob sie alles verstanden hat... (z. B. Falls du meinen Auftrag nicht vollständig verstanden hast, stelle mir zuerst Fragen bis du ganz sicher bist.) | Fragen? & Tools | |

22

Beim Metaprompting* steuern wir nicht die Antwort, sondern die Art, wie die KI zu ihren Antworten kommt.



*Der Begriff „Metaprompting“ ist nicht einheitlich definiert. In dieser Präsentation bezeichnet er die bewusste Steuerung der Denk- und Arbeitsweise der KI – nicht nur die Beschreibung des gewünschten Ergebnisses.

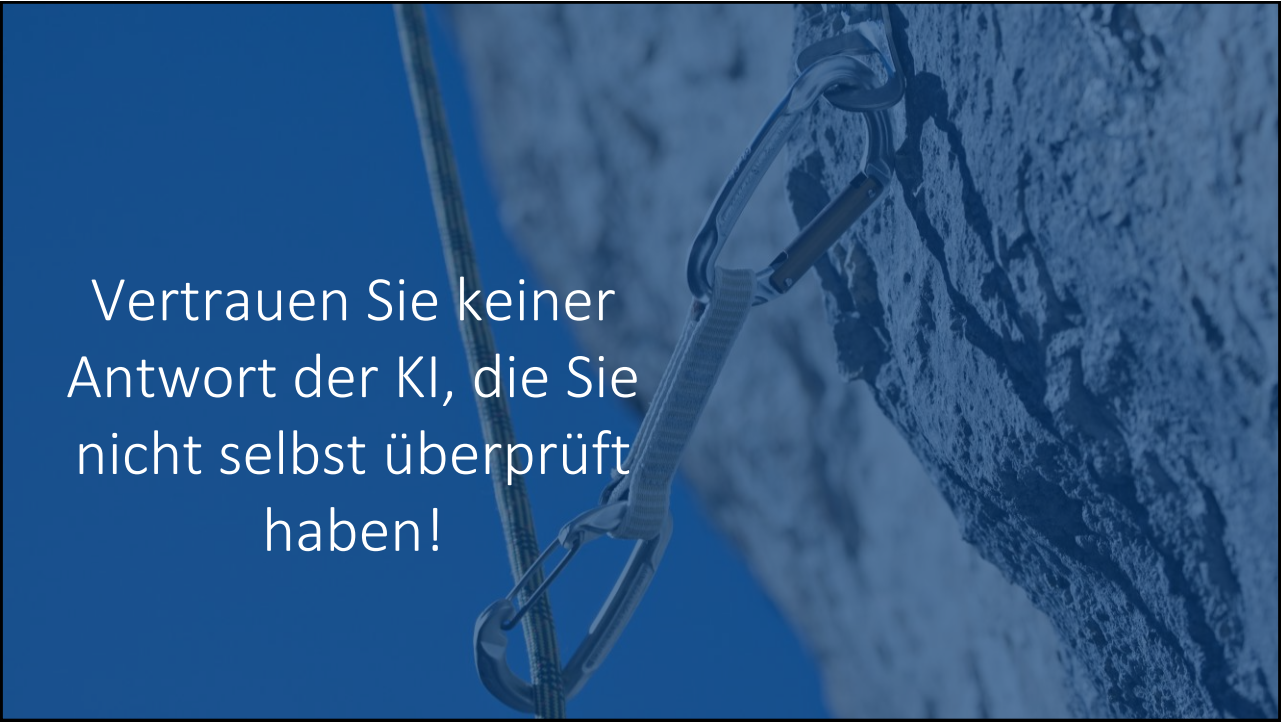
1 2 3 4 5 Metaprompting & Chain of Thought © Dr. S. Imboden Hes-so VALAIS WALLIS 23

Stärken und Grenzen von KI-Chatbots

Aufgabentyp	ChatGPT (OpenAI)	Claude (Anthropic)	Gemini (Google)	Perplexity	Copilot (Microsoft)	Grok (xAI)
Problemstrukturierung & heuristisches Denken (Thema ordnen, Denkpfade skizzieren)	●	●	●	●	●	●
Argumentative Tiefe & Kohärenz (stringent argumentieren, Pro & Contra)	●	●	●	●	●	●
Lange, nüchterne Fachtexte (Berichte, Studien, Konzepte)	●	●	●	●	●	●
Zusammenfassungen & Verdichtung (Kernaussagen aus Texten)	●	●	●	●	●	●
Systematische Recherche (Themenüberblick, Suchrichtungen)	●	●	●	●	●	●
Aktuelle Fakten & Quellen (Was gilt aktuell? Wo nachprüfen?)	●	●	●	●	●	●
Diskursive Einordnung & Perspektivenvielfalt (unterschiedliche Sichtweisen)	●	●	●	●	●	●
Kodieren & Programmieren (Code erklären, debuggen, strukturierende Analysen)	●	●	●	●	●	●
Tabellenkalkulation & Excel-Logik (Formeln, Datenlogik, Auswertungen)	●	●	●	●	●	●
Entscheidungsnahe Nutzung (Hochrisiko) (Medizin, Recht, heikle Entscheide)	●	●	●	●	●	●

● geeignet ● bedingt geeignet ● ungeeignet

- ChatGPT**
- Sehr stark in Strukturierung, Argumentation und erklärendem Denken
 - Breite Einsetzbarkeit über viele Aufgabentypen hinweg (Allrounder)
 - Risiko hoher Plausibilität bei inhaltlich unsicheren Ergebnissen
- Claude**
- Hohe argumentative Kohärenz und saubere Annahmexplizierung
 - Sehr gut für lange, neutrale Fachtexte
 - Teilweise zu vorsichtig und wenig explorativ
- Gemini**
- Gute Anbindung an aktuelle Informationen und Recherche
 - Solide Allround-Leistung für Standardaufgaben
 - Schwankende analytische Tiefe
- Perplexity**
- Sehr stark in Recherche, Exploration und Quellenorientierung
 - Gute Orientierung in neuen oder unbekannten Themenfeldern
 - Schwach bei Synthese und kohärenter Analyse
- Copilot**
- Sehr effizient in Office-, Excel- und Coding-Workflows
 - Niedrige Einstiegshürde in produktiven Alltagsaufgaben
 - Begrenzte konzeptionelle und argumentative Tiefe
- Grok**
- Pointierte, teils ungewöhnliche Perspektiven
 - Schnell bei diskursiven oder meinungsnahen Themen
 - Methodisch instabil und analytisch unzuverlässig



Vertrauen Sie keiner
Antwort der KI, die Sie
nicht selbst überprüft
haben!

25



Agenda

1. Kurze Begriffsklärung
2. Die vier wichtigsten Einschränkungen von Chatbots
3. Metaprompting & Chain of Thought
4. **Demo**
5. Fazit

26

Demo

1. Notebook LM
2. Metaprompting

27



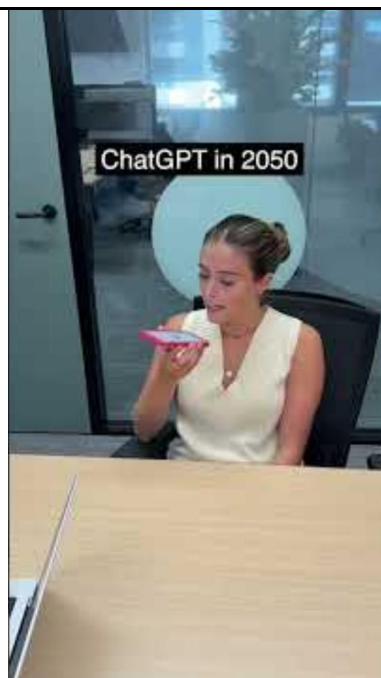
Agenda

1. Kurze Begriffsklärung
2. Die vier wichtigsten Einschränkungen von Chatbots
3. Metaprompting & Chain of Thought
4. Demo
5. Fazit

28

Fazit: Die Zukunft gehört nicht der KI – sondern den Menschen, die lernen, mit ihr zu denken

1. **Grenzen:** Generative KI halluziniert bisweilen (erfindet plausibel klingende, aber falsche Inhalte) und hat keine verlässliche Faktenbasis – daher Antworten nie blind übernehmen, sondern immer kritisch prüfen.
2. **TITO-Prinzip & Prompting:** „Trash in, Trash out“ – KI liefert nur so gute Resultate wie ihre Eingaben. Klare, präzise Prompts und hochwertige Daten werden zur neuen Schlüsselkompetenz.



<https://youtube.com/shorts/9VnYy-l5xsg?si=YPeOEGJl2Q8Ezrds>



Besten Dank für Ihre
Aufmerksamkeit



Hes·so VALAIS
WALLIS

Hochschule für Wirtschaft & Tourismus
Dr. Serge Imboden
Techno-Pôle 3
3960 Sierre
+41 27 606 90 72
+41 79 217 06 08
serge.imboden@hevs.ch
www.2iManagement.ch



