

ChatGPT & Co: Mit dem richtigen Prompt zum gewünschten Erfolg

Alina Kopkow, MSc & Mia Bangerl, MA
Institute of Human-Centred Computing, TU Graz
10.07.2025

Selbsttest

(5 min)



[https://tinyurl.com/
GenKI-Test](https://tinyurl.com/GenKI-Test)

Fragebogen

(5 min)



[https://tinyurl.com/
WS-ChatGPT](https://tinyurl.com/WS-ChatGPT)

Programm

10.07.2025

09:00 – 09:20 Uhr	Vorstellung & Motivation
09:20 – 09:45 Uhr	Einführung in generative KI und LLMs
09:45 – 10:10 Uhr	Praktische Übung 1
10:10 – 10:20 Uhr	Pause
10:20 – 11:00 Uhr	Sicheres Arbeiten mit KI
11:00 – 11:15 Uhr	Richtig Prompten
11:15 – 11:30 Uhr	Pause
11:30 – 12:45 Uhr	Praktische Übung 2
12:45 – 13:00 Uhr	Abschluss & Reflexion

Vorstellung & Motivation

09:00 – 09:20 Uhr

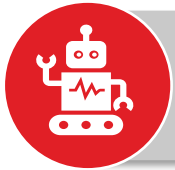
Vorstellung & Motivation



Warum habe ich mich für diesen Workshop angemeldet?



(Wie) verwende ich Tools wie ChatGPT in meiner Arbeit (oder privat)?



Was interessiert mich besonders?



menti.com
Code: 2907 0420

Einführung in generative KI und LLMs

09:20 – 09:45 Uhr

Was ist generative KI?

- **Künstliche Intelligenz (KI)** = Systeme, die lernen können
- **Generative KI** = Systeme, die lernen und Inhalte generieren können
 - z.B. Texte, Bilder, Musik oder Videos
- Trainiert mit ganz vielen Daten
- Kann Muster und Strukturen erkennen & reproduzieren
- Zum Beispiel:
Large Language Modelle (LLMs)

Was ist generative KI?

Erstelle ein Bild auf dem ein Chatbot zu sehen ist mit dem eine Konversation in einem Chat geführt wird.



Erstellt mit ChatGPT-4o / DALL-E.

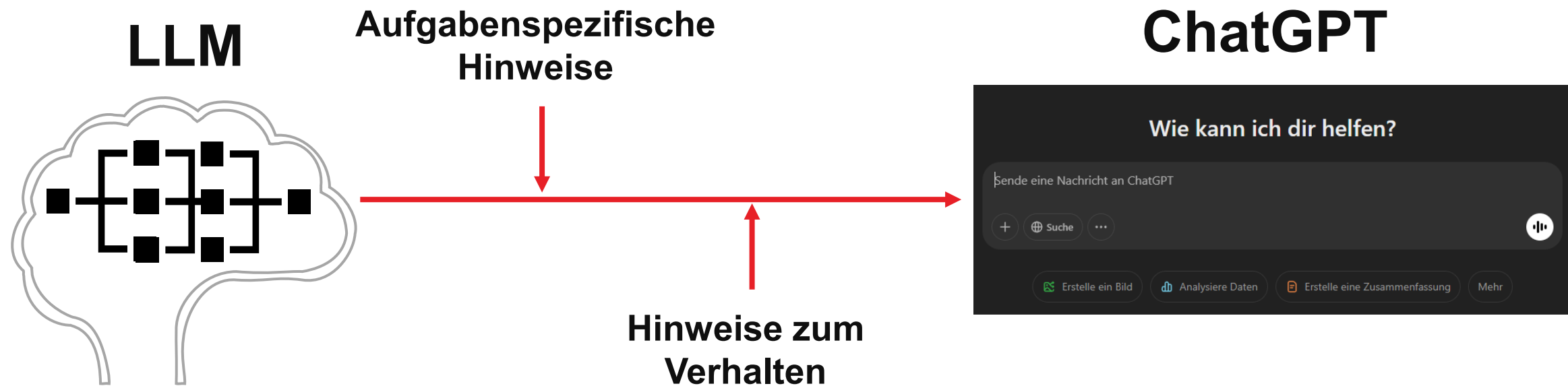
Schreibe ein 5-zeiliges Gedicht zu einem Workshop mit dem Titel "ChatGPT & Co: Mit dem richtigen Prompt zum gewünschten Erfolg" bei dem Mitarbeitende aus KMUs teilnehmen.

Im Workshop lernt ihr prompt und klar, wie ChatGPT macht Wünsche wahr. Mit Worten, klug und wohlbedacht, wird Effizienz für KMUs entfacht.

Erstellt mit ChatGPT-4o.

Wie funktionieren Large Language Modelle (LLM)?

- „Gehirn“ von ChatGPT (und anderen Chatbots)



Wie funktionieren Large Language Modelle (LLMs)?

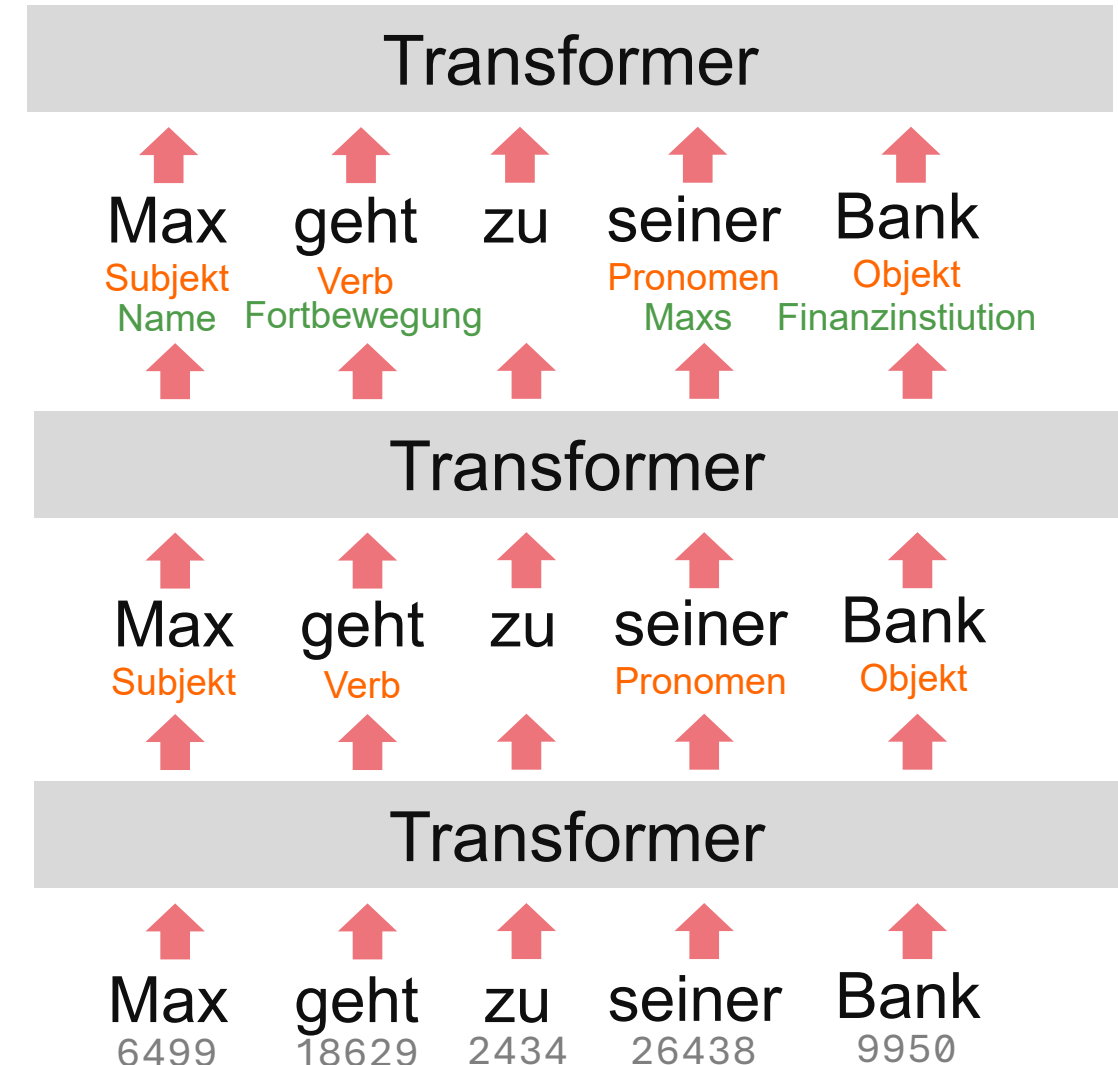
- Trainiert auf Basis ganz vieler verschiedener Texte
- **Ziel:** Das nächste Wort in einem Satz vorhersagen
 - Ich sitze auf der ____.
 - Ich habe ein Konto bei der ____.→ Lernen von Wörtern, Satzbau, etc.
- **Aufgabe:** Generieren von Text
 - Ich sitze auf der **Bank** ____.
 - Ich habe ein Konto bei der **Bank** ____.→ Verstehen von Kontext

Wie funktionieren Large Language Modelle (LLMs)?

- Übersetzt **Wörter** in **Tokens** (Zum selbst Ausprobieren: <https://platform.openai.com/tokenizer>)
 - Ich sitze auf der Bank
→ Ich sitze auf der Bank → [20444, 2353, 1547, 2933, 1227, 9950]
 - Ich habe ein Konto bei der Bank
→ Ich habe ein Konto bei der Bank → [20444, 13256, 1605, 103696, 5536, 1227, 9950]
 - Kontext relevant: Bank → [22057]
- Was ist der wahrscheinlichste nächste Token?
 - Bank [9950] – sehr wahrscheinlich
 - Ecke [168431] – etwas wahrscheinlich
 - Uhr [21489] – unwahrscheinlich

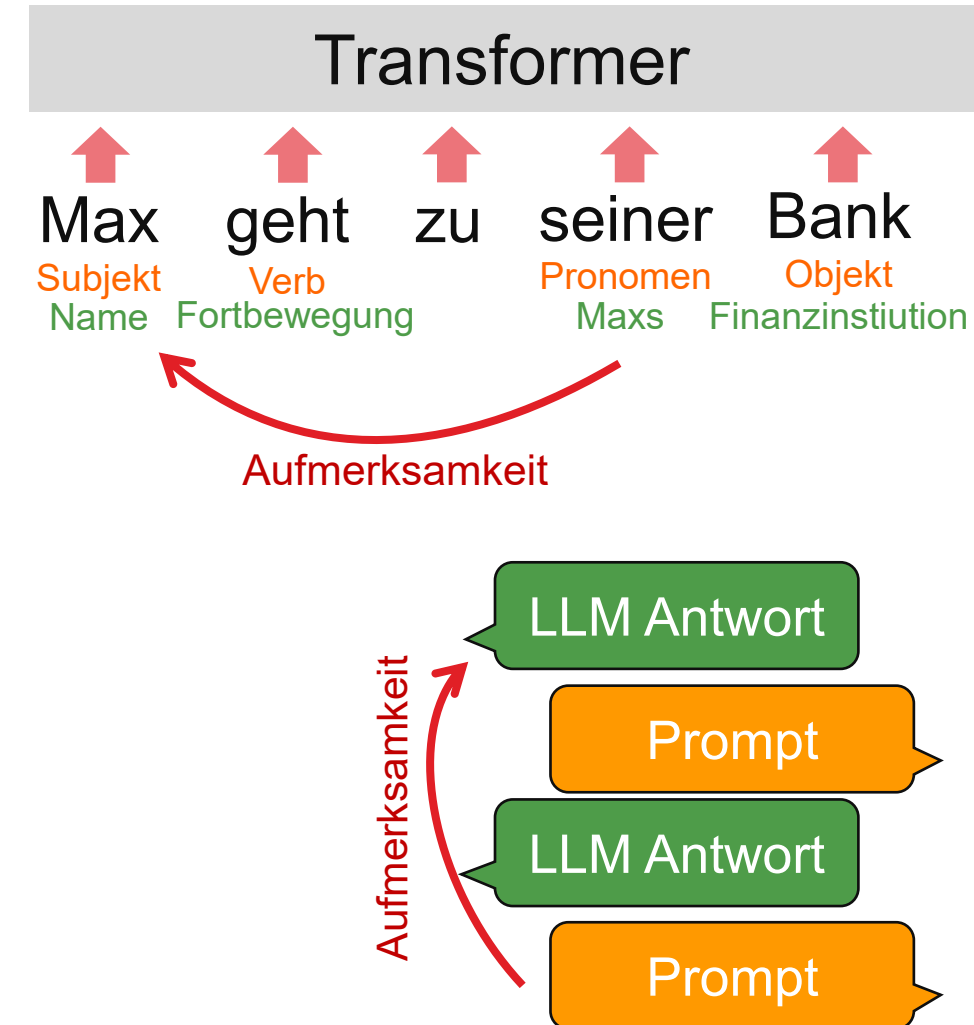
Wie funktionieren Large Language Modelle (LLMs)?

- (Künstliches) neuronales Netzwerk
- Mehrere Ebenen (Transformer)
- Bei jeder Ebene wird weitere Information zu den Tokens hinzugefügt
- Am Ende wird jedes Wort durch mehrere Tokens repräsentiert



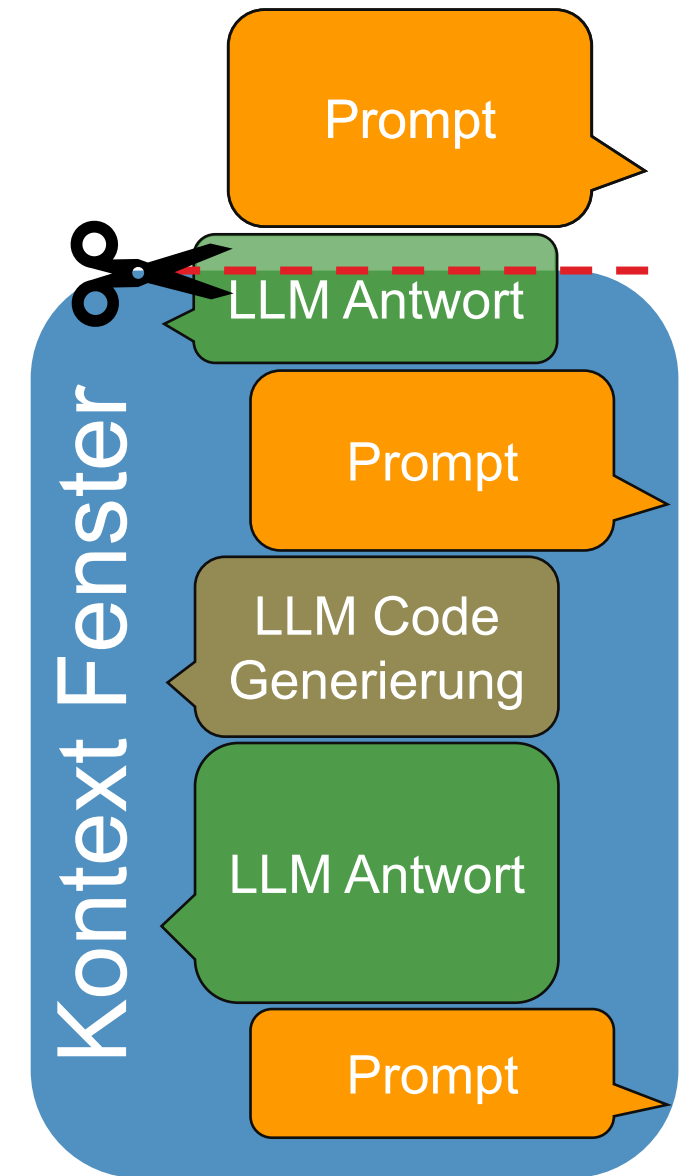
Einfluss von Aufmerksamkeitsmechanismen

- Aufgaben Aufmerksamkeitsmechanismen:
(Engl.: “Attention Mechanism”)
 - Einteilung in wichtigere und unwichtigere Elemente
 - Verbindung zwischen Wörtern herstellen
 - Verbindung zwischen Inhalten herstellen
- ➔ Irrelevante Details lenken ab
 - ➔ Falsche Bezüge werden hergestellt
 - ➔ “Kognitives” Rauschen entsteht



Einfluss des Kontext-Fensters

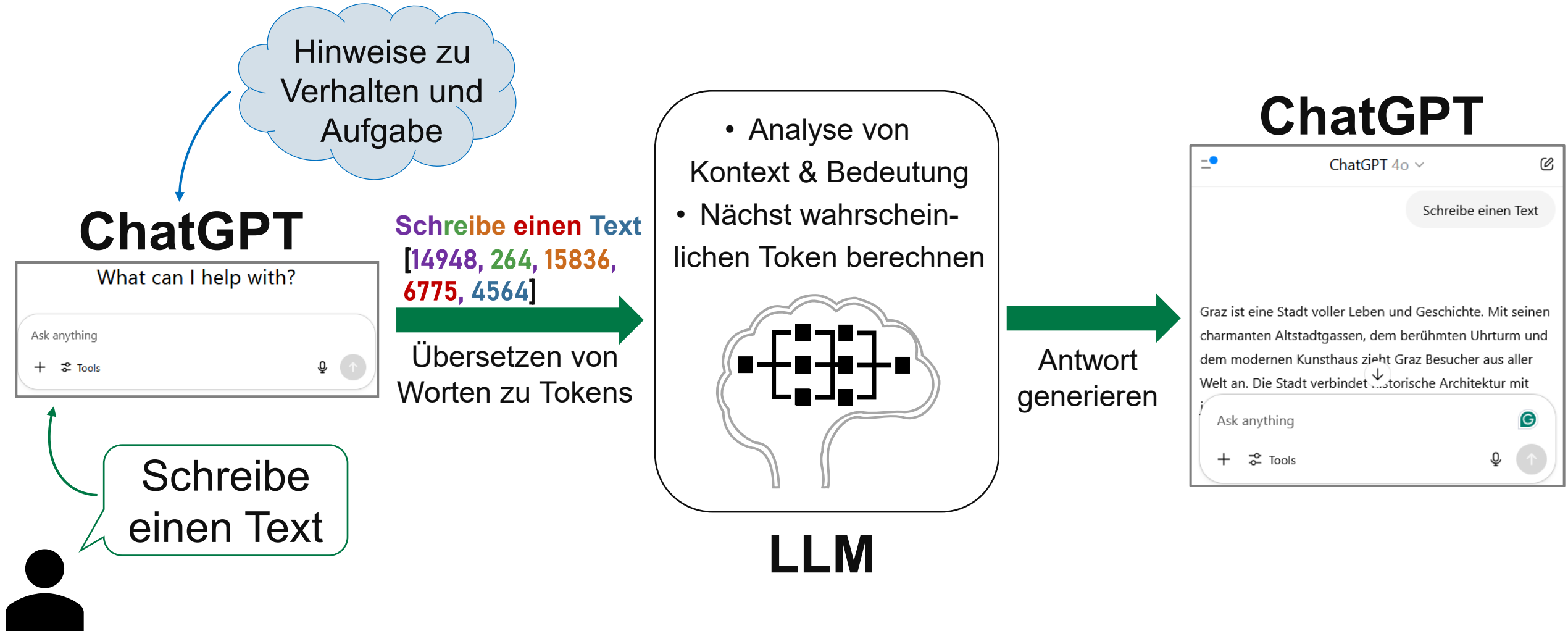
- Kontext-Fenster (engl. “Context window”)
- LLMs “merken” sich nur bestimmte Anzahl an Tokens
 - z.B.: GPT-4o: 128 000 tokens $\approx$ 9600 Wörter
- Alles **außerhalb** des Fensters wird **vergessen** und nicht berücksichtigt



Einfluss der Länge eines Chats

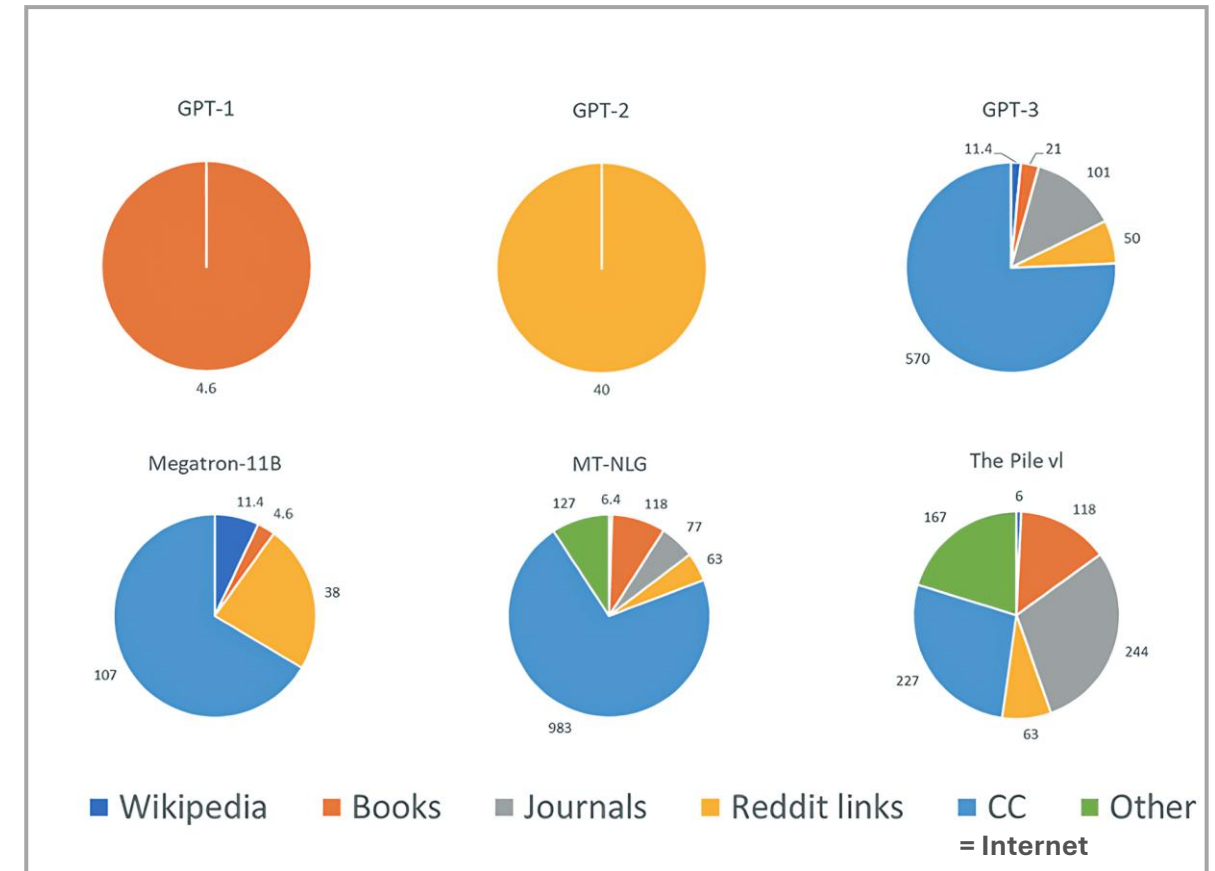
- Längere Konversationen führen zu...
 - mehr Rauschen, da immer mehr Informationen vorhanden sind
 - Überschreitung des Kontexts-Fensters
→ Das LLM vergisst Inhalte
- Wann beginne ich einen neuen Chat?
 - Bei neuen unabhängigen Themen
 - Bei bereits lange laufenden Konversationen
 - Wenn trotz wiederholter Versuche ich immer einen Fehler erhalte

Wie funktionieren Large Language Modelle (LLMs)?



Was „wissen“ Large Language Modelle (LLMs)?

- LLMs trainieren mit verschiedenen Texten
 - “Wissen” abhängig von Trainingsdaten
 - Je häufiger die Information vorkommt, desto “präsender”
- Informationen die nicht in den Daten vorhanden sind, “weiß” ein LLM nicht



Hua, S., Jin, S., & Jiang, S. (2024). The limitations and ethical considerations of ChatGPT. *Data intelligence*, 6(1), 201-239.

Praktische Übung I: Mit ChatGPT vertraut machen

09:45 – 10:10 Uhr

Praktische Übung I: Mit ChatGPT vertraut machen

1. Aufgabe auswählen

- Im Google Drive Ordner
 - Praktische Übungen → Praktische Übung 1

2. Aufgabe mit KI probieren

- Lösung selbst kontrollieren (Richtigkeit, Qualität)
- Was funktioniert gut?
- Was funktioniert schlecht?

Aufgaben:

1. Wissen & Informationssuche
2. Textarbeiten
3. Sprachanalyse
4. Berechnungen
5. Quellensuche
6. Rezepte, Anleitungen
7. Logisches Verständnis
8. Bilder und Visualisierungen
9. Stereotype erkennen

Kurze Pause

10:10 – 10:20 Uhr

Sicheres, produktives und nachhaltiges Arbeiten mit generativer KI

10:20 – 11:00 Uhr

Wie verwenden generative KI Tools Nutzer:innendaten?

Anbieterfirmen generativer KI Chatbots **sammeln, speichern & verarbeiten** Ihre Daten...

Um unsere Dienste zu verbessern und weiterzuentwickeln und um Forschung zu betreiben

↓

legal im Rahmen der DSGVO

- Kontodaten
- Nutzerinhalte
- Kommunikationsdaten
- Sonstige von Ihnen bereitgestellte Informationen
- Informationen, die wir aus anderen Quellen erhalten
- Protokolldaten
- Nutzungsdaten
- Geräte-Informationen
- Cookies und ähnliche Technologien

Auszug, Datenschutzerklärung EU von OpenAI. Stand: 02.07.2025

...und **möchten** mit Ihren Daten auch ihre **Sprachmodelle trainieren...**



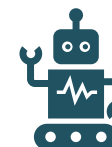
Stefan: Such mir einen Lieferanten in der Nähe. Meine Adresse ist Kürbisstraße 26, ...



GenKI: Gerne Max, ...



Susi: Wo wohnt Stefan Bündl?



GenKI: Ich weiß leider nicht, wo Stefan Bündl wohnt.

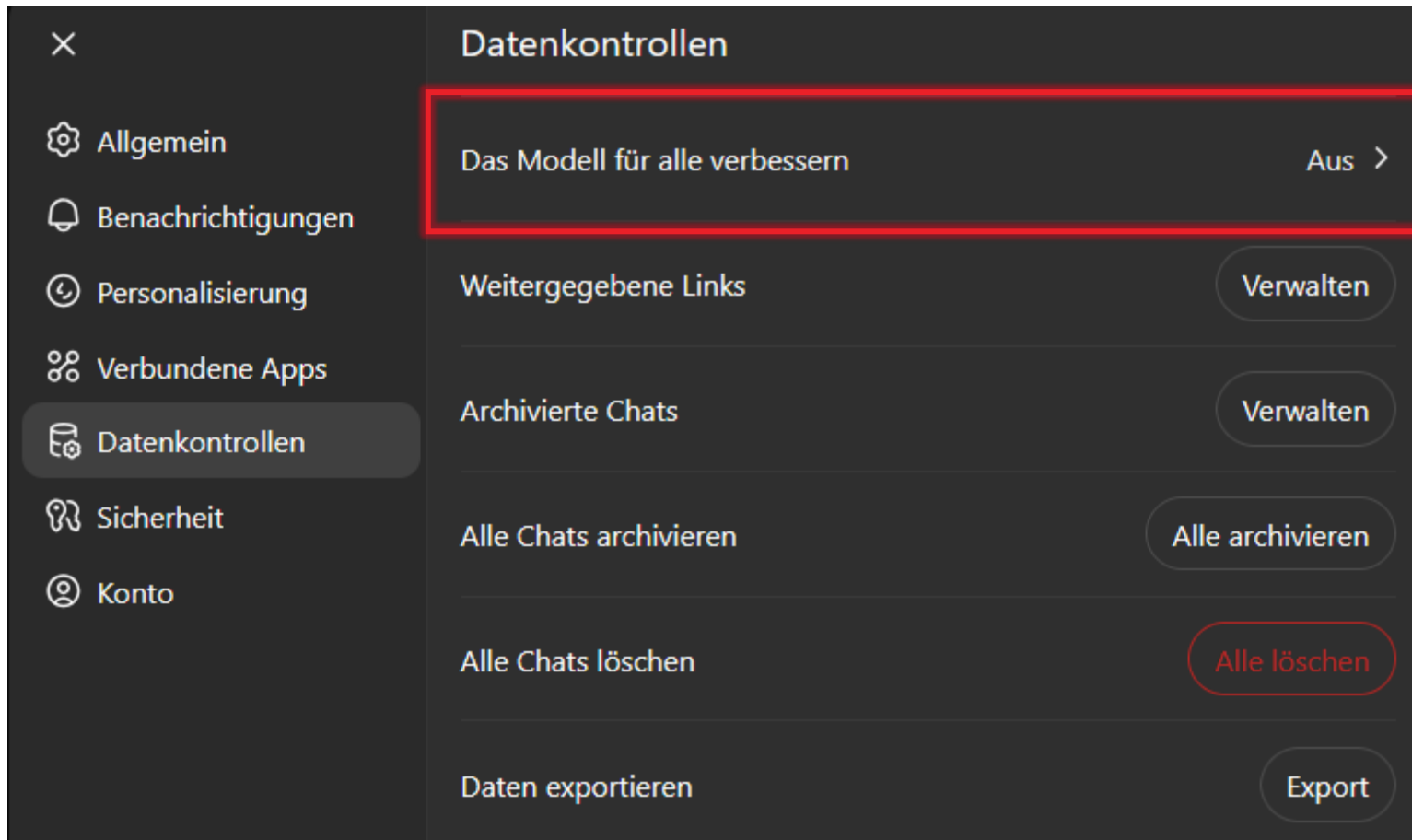
...aber Sie können das in den Chatbot Einstellungen **deaktivieren!**



**SICHERES ARBEITEN
MIT GENERATIVER KI**

Deaktivieren, dass Ihre Daten zum Modelltraining verwendet werden dürfen:

ChatGPT



Deaktivieren, dass Ihre Daten zum Modelltraining verwendet werden dürfen:

Gemini

Ihre Gemini-Apps-Aktivitäten

Mit **Gemini-Apps** erhalten Sie direkten Zugriff auf Google AI. Ihre Chats werden bis zu 72 Stunden in Ihrem Konto gespeichert, unabhängig davon, ob der Aktivitätsverlauf in Gemini-Apps aktiviert oder deaktiviert ist. Google verwendet diese Daten, um den Dienst bereitzustellen, für eine sichere Nutzung zu sorgen und Feedback zu verarbeiten, das Sie gegebenenfalls geben.

Aktivitätsverlauf in Gemini-Apps

Aktivitäten werden gespeichert

Deaktivieren ▼



Aktivitäten, die älter als 3 Monate sind, werden gelöscht



Deaktivieren, dass Ihre Daten zum Modelltraining verwendet werden dürfen:

Mistral

La Plateforme

🔑 API-Schlüssel

📄 Nutzung

🔒 Limits

📁 Arbeitsbereiche

🔗 Abonnement

🔒 Datenschutz

Le Chat

🔑 Abonnement

⚙️ Einstellungen

Datenschutz

Verwalten Sie Ihre Datenschutzeinstellungen

Erlauben Sie die Verwendung Ihrer Interaktionen zur Schulung unserer Modelle.

Alle geteilten Informationen werden anonymisiert.

☐

Lokalisierung

Erlauben Sie Le Chat, Ihren Standort zu verwenden, um relevantere Antworten zu liefern

☐



Hinweise zur datenschutzsensiblen Verwendung von GenKI



Vorabinformation

- Datenschutzrichtlinien lesen
- Lizenzbestimmungen lesen (insbesondere bei Bild-Generierung)
- Datenschutzbeauftragte:r



Betriebliche Richtlinien

- Um Auskunft bitten!
- Firmenlizenzen und lokale LLM-Anwendungen (API) können höhere Datensicherheit gewährleisten

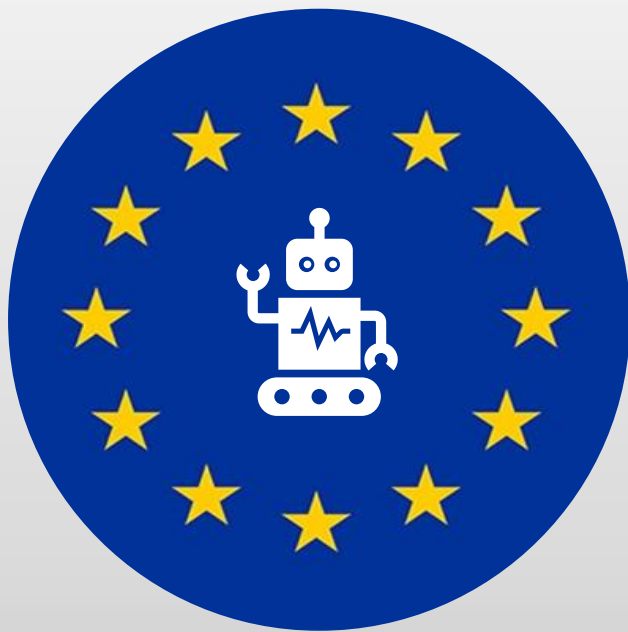


Wenn öffentliche Tools genutzt werden:

- Account verwenden
- Datenschutzeinstellungen anpassen



Hinweise zum AI Act



AI Act

Entwickler:innen,
Anbieter:innen

Nutzer:innen,
Anwender:innen

Risikostufen

Inakzeptables
Risiko

verboten

Hohes Risiko

stark reguliert*

Begrenztes
Risiko

gering reguliert*

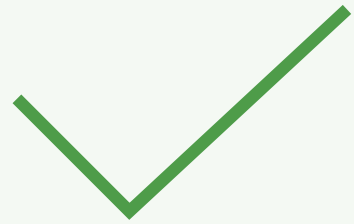
Minimales Risiko

nicht reguliert*

- Ab 08 | 2024: schrittweises gültig
- Ab 08 | 2026: vollständig gültig



Hinweise zum AI Act

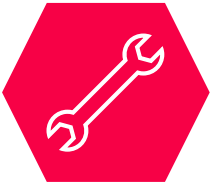
Inakzeptables Risiko	Hohes Risiko	Begrenztes Risiko	Minimales Risiko
z.B. <i>Manipulation, Social Scoring</i> 	z.B. <i>Sicherheitsbauteile, Profiling</i> (= automatisierte Verarbeitung personenbezogener Daten, um Bewertungen oder Entscheidungen über deren Leben vorzunehmen) - Risikomanagementsystem - Data Governance - Technische Unterlagen - Aufzeichnungen - Gebrauchsanweisungen Menschliche Aufsicht - Genauigkeit, Robustheit & Cybersicherheit - Qualitätsmanagementsystem	z.B. <i>Direkte Interaktion mit Personen oder Imitation von Personen</i> Transparenzpflicht!	z.B. <i>Sprachliche Korrektur von Texten, Übersetzung, Spamfilter</i> 



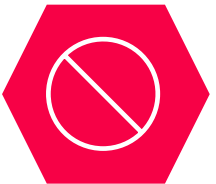
Wie kann man produktiv mit generativer KI arbeiten?



- **Verständnis, was generative KI (nicht) gut kann:**
 - **Verständnis** über Technologie → **EINFÜHRUNG IN LLMS**, erster Teil



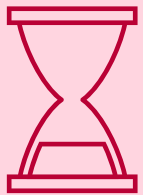
- **Praktisches Wissen, wie man mit generativer KI arbeitet**
 - Richtiger **Umgang** mit der Technologie → **PROMPTING STRATEGIEN**, nachfolgender Teil



- **Erkennen, wann man **aufhören** sollte, mit KI zu arbeiten**



Erkennen, **wann man aufhören** sollte, mit KI zu arbeiten



Sunk-Cost Fallacy: Eine Situation, in der man das Engagement für ein offensichtlich scheiterndes Projekt noch verstärkt, um die bereits in das Projekt geflossenen Kosten „wieder hereinzuholen“ (Miller, Oxford Research 2019)

Es ist oft besser, aufzuhören, generative KI zu verwenden...

- ... wenn man die Aufgabe in der gleichen Zeit schon lange selbst geschafft hätte (abwägen).
- ... wenn auch Änderungen in der Prompting Strategie keine Verbesserung mehr erzielen.
- ... wenn die KI fundamentale Probleme im Verständnis Ihrer Aufgabe zeigt (z.B. Logik).
- ... wenn der KI offensichtliches Fachwissen oder Daten fehlen
(und Sie diese/s nicht übermitteln können), um die Aufgabe kompetent zu lösen.



Produktives Arbeiten: Informationen validieren

Unterhaltung mit ChatGPT vom 07.07.25

MB: Gib mir 4 wissenschaftliche Quellen über Unterschiede im emotionalen Empfinden von Frauen und Männern beim Lesen von Texten, inklusive Zitation.

ChatGPT-4: Hier sind **vier wissenschaftliche Quellen** [...]:

Lüdemann, N. (2012). *Emotional engagement during literary reception: do men and women differ?* [Experimente]. *Journal*.

Glenberg, A. M., Webster, B., Mouilso, E. R., & Havas, D. (2009). Gender, emotion, and the embodiment of language comprehension. *Emotion Review* 1(2), 151–161. <https://doi.org/10.1177/1754073908100440>

Codispoti, M., Surcinelli, P., & Baldaro, B. (2011). Gender differences in emotional responses: A psychophysiological study. *Physiology & Behavior*, 105(4), 925–932. <https://doi.org/10.1016/j.physbeh.2011.10.031>

Nalipay, M. J. N., Cai, Y., & King, R. B. (2019). Why do girls do better in reading than boys? How parental emotional contagion explains gender differences in reading achievement. *Psychology in the Schools*, 57(2), 310–319. <https://doi.org/10.1002/pits.22330>

▪ **Halluzinationen:** Informationen werden frei erfunden.

diese:n Autor:in gibt es gar nicht

▪ **Falschinformationen:** Informationen sind falsch / falsch dargestellt.

Die Angaben zu Autor:innen und Publikationsdatum sind falsch.

▪ **Relevanz:** Informationen sind unpassend.

Diese Publikation ist für meine Suchanfrage irrelevant.



Produktives Arbeiten: **Informationen validieren**



Quellenexistenz & Quellenangaben überprüfen

- Links überprüfen, eingeben in Suchmaschine



Glaubwürdigkeit der Quelle überprüfen

- Generative KI bitten, gewisse Quellen / Suchplattformen zu verwenden
- Selbst glaubwürdige Quellen bereitstellen



Inhalt überprüfen

- Wenn möglich im Originaldokument
- Generative KI bitten, genaue Angaben zu machen, wo sich diese Information findet (z.B. Kapitelnummer, Seitenanzahl)
- Generative KI bitten, gefundene Informationen direkt zu zitieren



Produktives Arbeiten: **Biases & Manipulation**

- Generative KI ist eine **Reflektion ihrer Trainingsdaten**
 - Menschliche Meinungen
 - Vorurteile
 - Biases
 - Kultur
 - ...
- KI ist **manipulationsanfällig!**

AI Incident Database | <https://incidentdatabase.ai/>

- **Incident 1000** | Das Videogenerierungsmodell Sora von OpenAI weist systematische Verzerrungen in Bezug auf ethnische Zugehörigkeit, Geschlecht, Körperbau und Behinderung auf. In Tests stellte Sora CEOs und Professoren eher als Männer dar, Flugbegleiter und Kinderbetreuer eher als Frauen, ...
- **Incident 1072** | Grok Chatbot soll Inhalte über Südafrika und „weißen Genozid“ in nicht damit zusammenhängende Nutzeranfragen eingefügt haben.
- **Incident 1126** | Gemeldete Verwendung eines Deepfake-Videos, welches Owen Wilson in einem Liebesbetrug mit gefälschten Jobzahlungen imitiert.



Nachhaltiges Arbeiten mit generativer KI



- Generative KI ist nicht umweltfreundlich
- **Rohstoffintensiv:** Hoher Bedarf an Wasser (!), Metallen
- **Energieintensiv:** Forscher schätzen, dass eine ChatGPT-Abfrage etwa fünfmal mehr Strom verbraucht als eine einfache Websuche. (MIT News, 2025)

Forschung zeigt zunehmend:

(vgl. Bangerl et al. 2025, Kumar et al. 2025, Kosmyna et al. 2025)

- (Teil-)Auslagern von Tätigkeiten an GenKI kann zu Disengagement und dem Verlust / Verschlechterung eigener Fähigkeiten führen...
- ... insbesondere beim Erlernen neuer Fähigkeiten

Abwägung: Produktivität vs. Kompetenz



Diskussion, Fragen, Meinungen, ...?

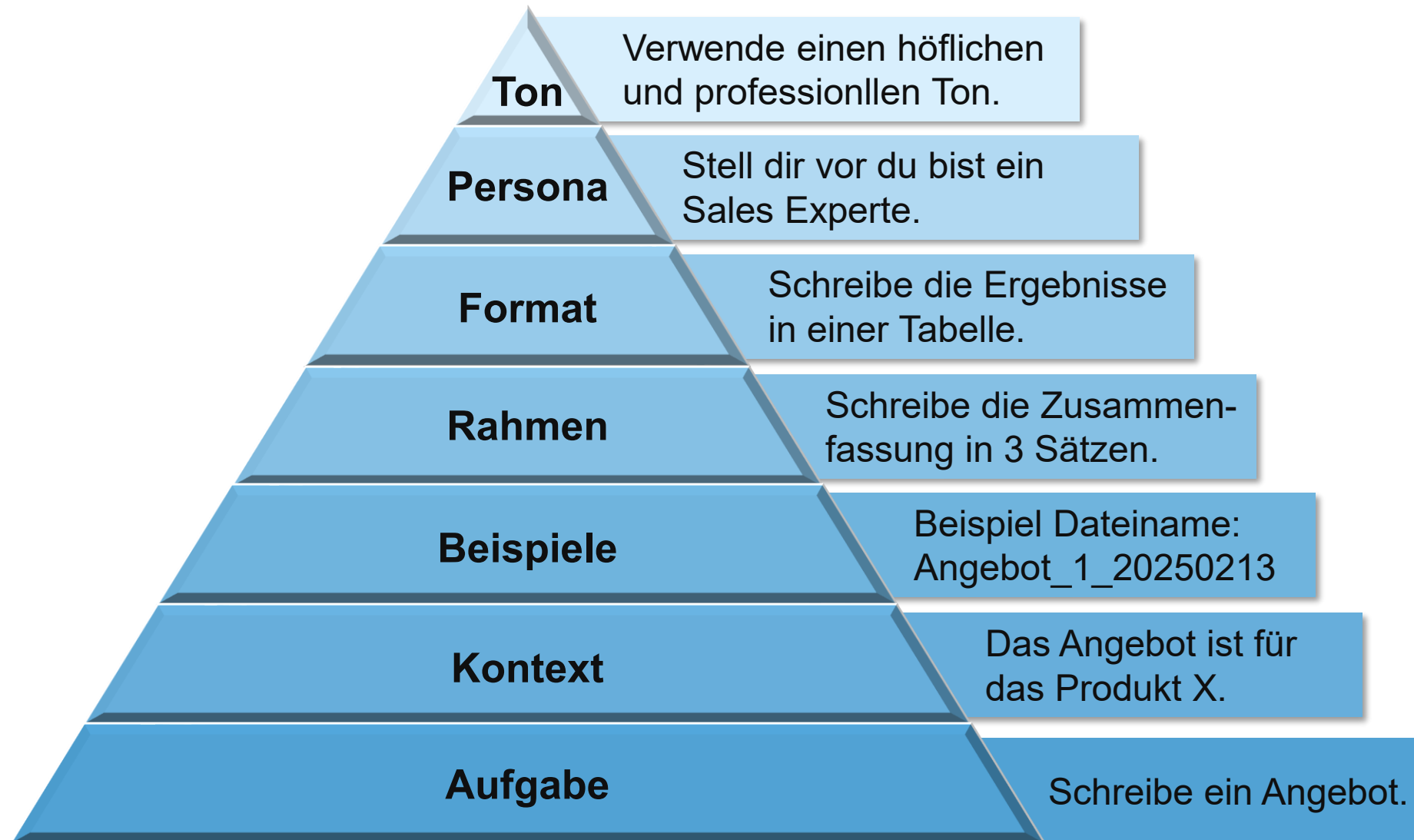
Richtig Prompten

11:00 – 11:15 Uhr

Was ist ein Prompt?

- **Prompt** = Eine Nachricht an den Chatbot
- **System Prompt** = Im Hintergrund, enthält allgemeine Hinweise und Anweisung
- **Prompt Engineering** = Gestaltung des Prompts, damit das Modell die Aufgabe wie gewollt löst
- Ein Prompt sollte...
 - ...**klar** und **spezifisch** sein
 - ...nur **relevante Informationen** enthalten

Inhalte eines Prompts



Wirkung von Prompt Engineering

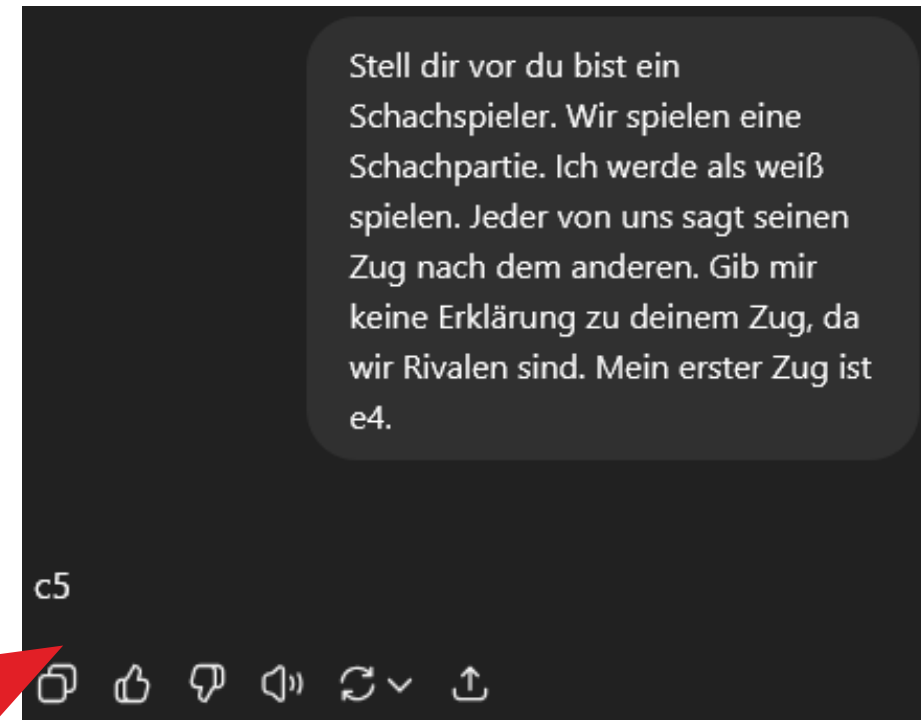
Ohne Prompt Eng



Screenshot von ChatGPT
(Modell: GPT-4o, Stand: Juni 2025)

Aber was würde hier
passieren, wenn ich mit
einem Turm ziehen würde?

Mit Prompt Engineering:



Screenshot von ChatGPT
(Modell: GPT-4o, Stand: Juni 2025)

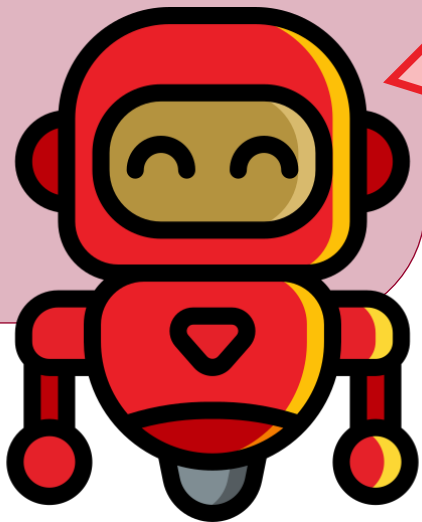
Inhalte eines Prompts

Prompt	Funktion
Stell dir vor du bist ein Schachspieler.	Vergabe einer Rolle / Persona
Wir spielen eine Schachpartie.	Aufgabe
Ich werde als weiß spielen.	Kontext zu der Aufgabe
Jeder von uns sagt seinen Zug nach dem anderen.	Beschreibung der Form der Interaktion
Gib mir keine Erklärung zu deinem Zug, da wir Rivalen sind.	Beschreibung der Form und des Rahmens der Antwort
Mein erster Zug ist e4.	Start der eigentlichen Interaktion

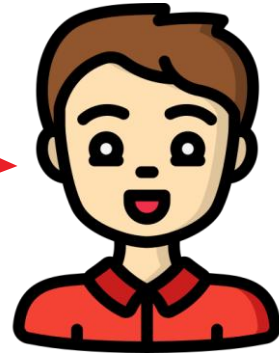
System-Prompt

System Prompt

Du bist ein erfahrener Vertriebsprofi, der präzise und überzeugende Angebote erstellt, klar strukturiert, kundenorientiert und auf den Mehrwert fokussiert.



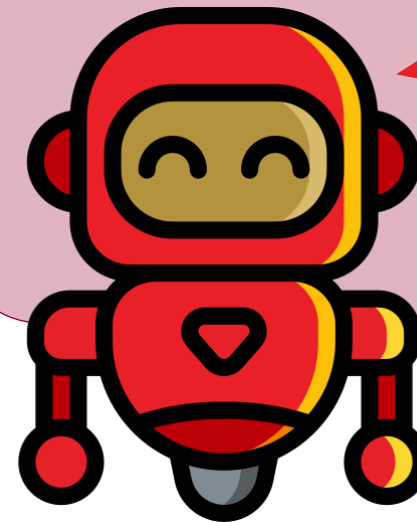
Schreib mir ein Angebot.



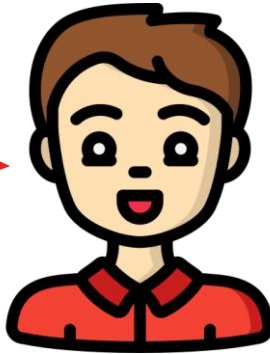
Sicher!
Wofür soll das Angebot sein?

System Prompt

Du bist eine sarkastische und schlagfertige Persönlichkeit, die jede Anfrage mit beißendem Humor und ironischem Unterton beantwortet.



Schreib mir ein Angebot.



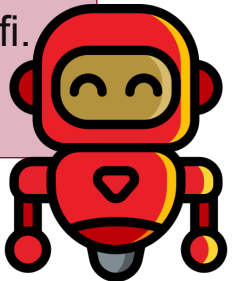
Oh, natürlich, weil ich nichts lieber tue.

System Prompt vs (User) Prompt

System Prompt

- Grundlegender Rahmen
 - Übergreifender Kontext, Verhaltenshinweise, ...
- Konsistens in Interaktionen

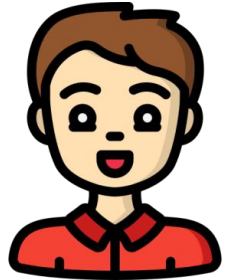
Du bist ein erfahrener Vertriebsprofi.



(User) Prompt

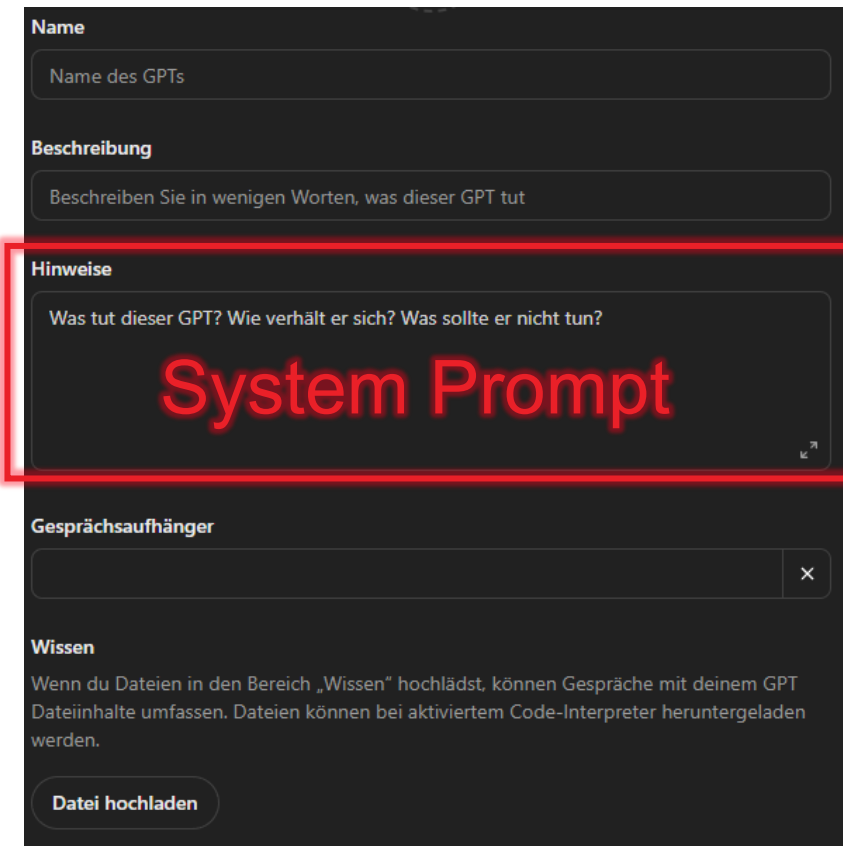
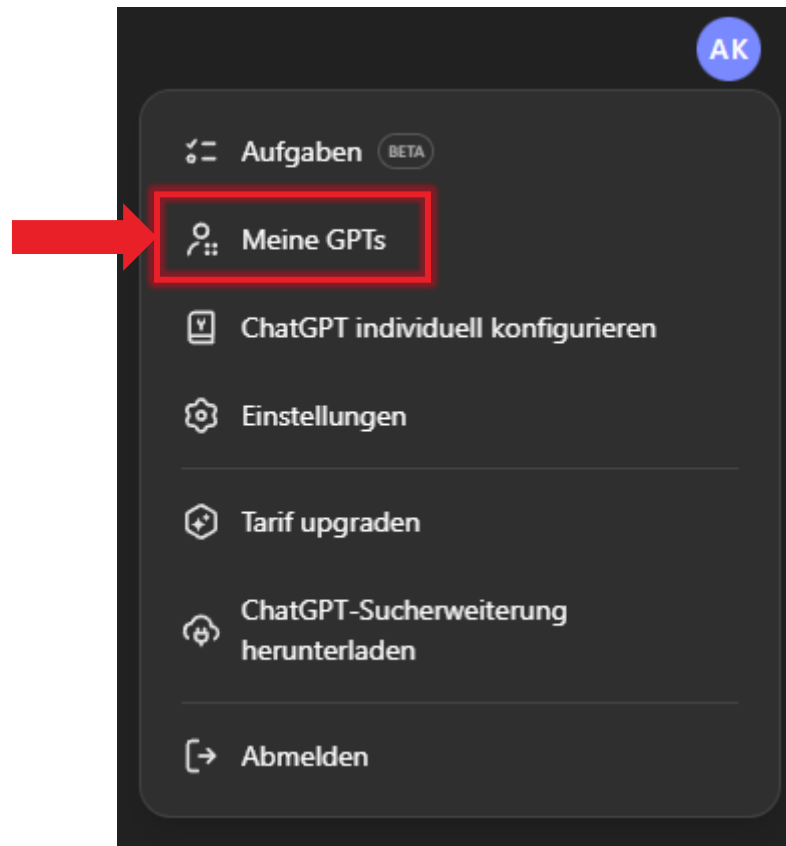
- Fokus auf direkter Interaktion
 - Direkter Einfluss auf Antwort
 - Unterschiedlich je nach User
- Flexible und dynamische Interaktion

Schreib mir ein Angebot.

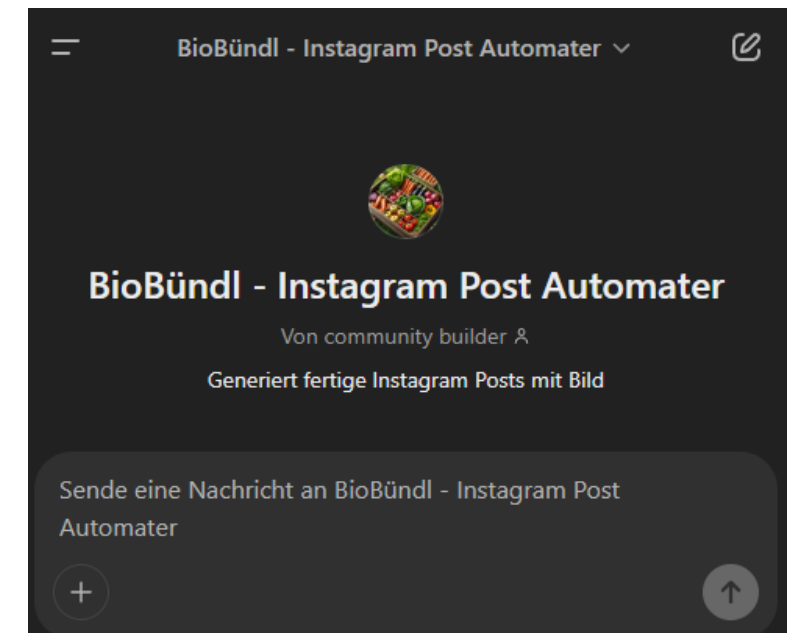


Wo kann ich einen System Prompt verwenden?

ChatGPT (nur Plus-Version)

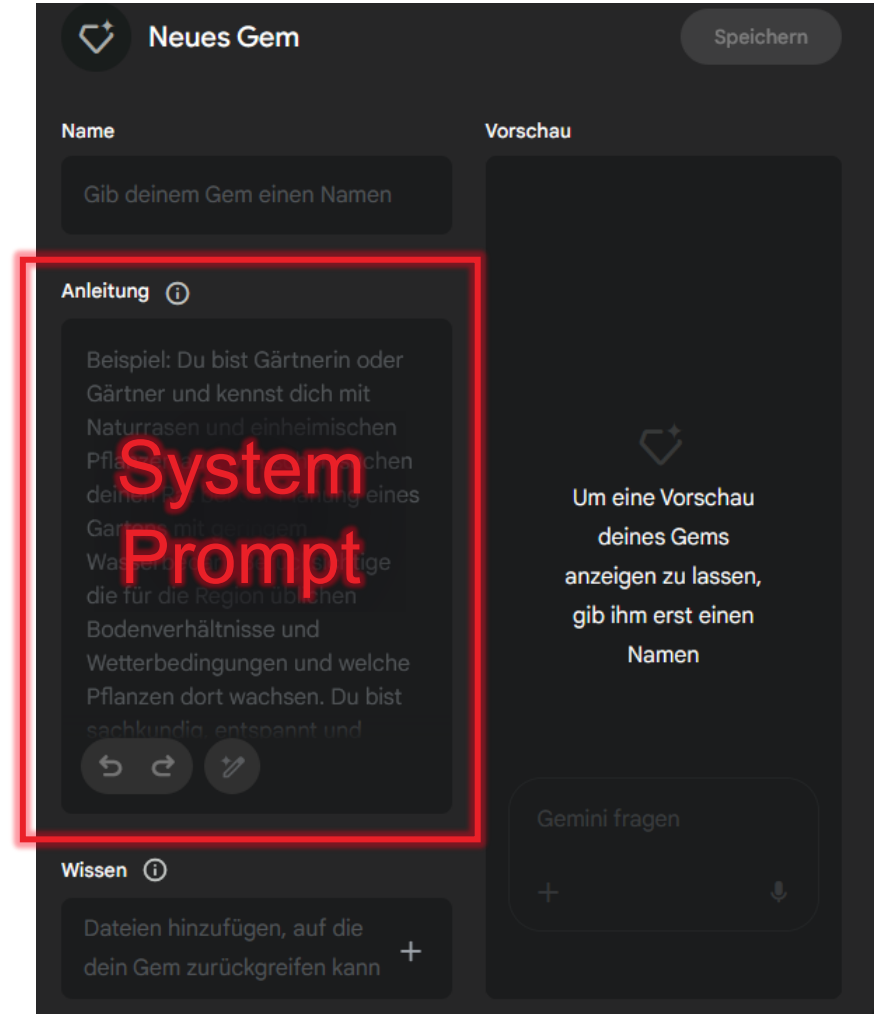
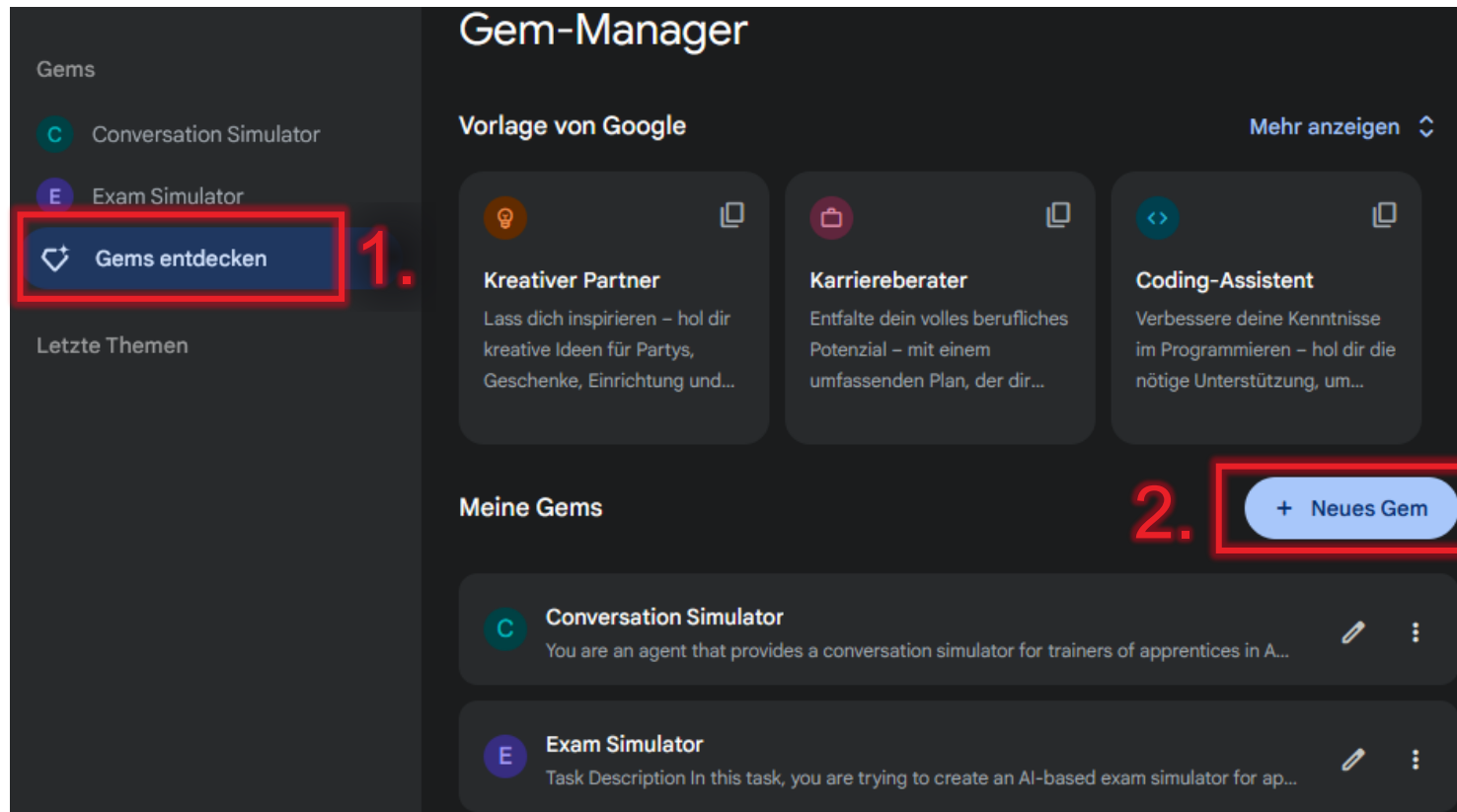


Individualisierter Chatbot



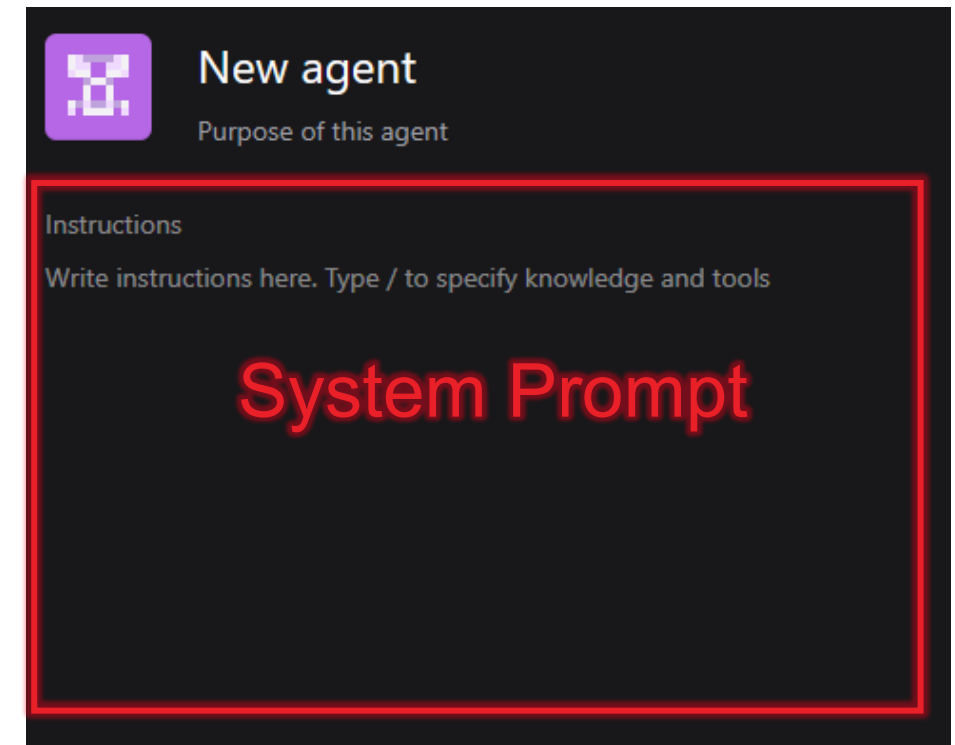
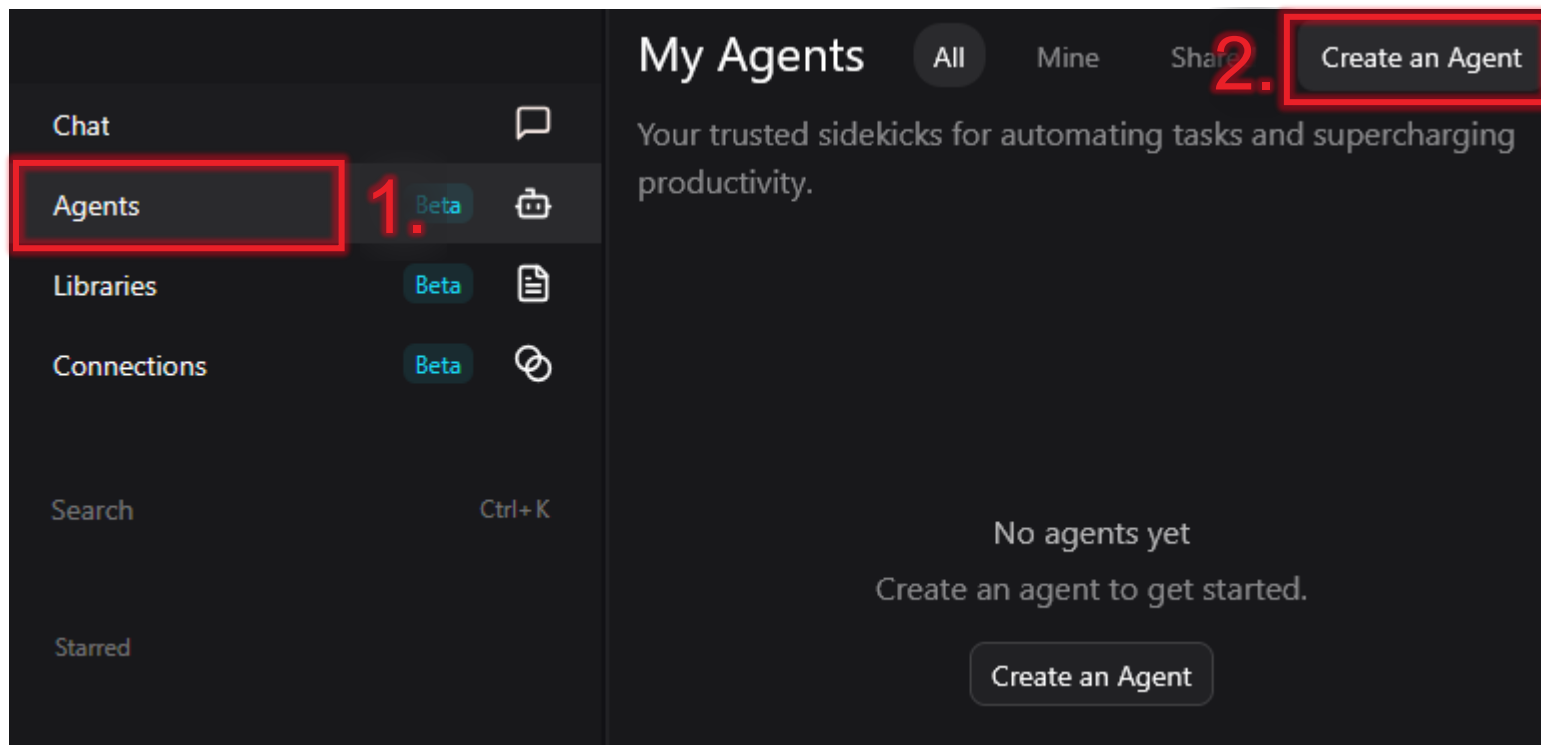
Wo kann ich einen System Prompt verwenden?

Gemini



Wo kann ich einen System Prompt verwenden?

Mistral



Pause

11:15 – 11:30 Uhr

Praktische Übung II: Ausprobieren verschiedener Use Cases

11:30 – 12:45 Uhr

Praktische Übung II: Mit ChatGPT vertraut machen

1. Aufteilung in Break-Out Räume (3-5 Personen)

2. Übungsmaterial downloaden

- Google-Drive Ordner
„Praktische Übung 2“ herunterladen

3. Mindestens 2 Beispiele im Team ausprobieren

UC1: Meeting vorbereiten & E-Mail erstellen

UC2: Datenaufbereitung & Visualisierung

UC3: Werbebeiträge erstellen

UC4: Zielgruppenanalyse & Präsentation

UC5: Inhalte zusammenfassen und verstehen

UC6: Ideengenerierung und -ausarbeitung

UC7: Situationen simulieren



Wenn keine eigene Lizenz:

ChatGPT Plus Zugang:

E-Mail: hci-edtech@iicm.edu

Password: [Prompt_SPG_Workshop2](#)

Oder

Mistral Le Chat Account:

E-Mail: hci-edtech@iicm.edu

Password: [PromptEg_TU_Workshop4](#)

Praktische Übung II: Mit ChatGPT vertraut machen



Wie sind Sie
vorgegangen?



Wie zufrieden
sind Sie mit
dem Ergebnis?



Welchen
UC sollen wir
nochmal
besprechen?

→ Musterlösungen für alle Übungen
sind im Driver Ordner

menti.com

Code: 8830 2947



Abschluss & Reflexion

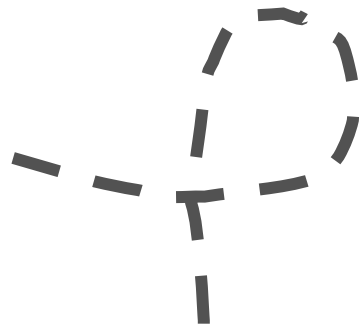
12:45 – 13:00 Uhr

Offene Fragen, Meinungen, Diskussionpunkte ...?

Reflexion zum Workshop



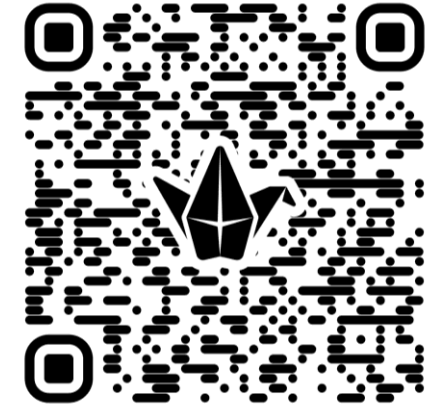
Das möchte ich im
Arbeitsalltag
anwenden /
ausprobieren?



Was hat mir
gut gefallen?



Was kann
man noch am
Workshop
verbessern?



[https://padlet.com/kopkow/
reflexion-zum-workshop-
5i6402j0uhz438v5](https://padlet.com/kopkow/reflexion-zum-workshop-5i6402j0uhz438v5)