

Pen and Paper – AI Training

Beschreibung:

Große Sprachmodelle (LLMs) werden vielfältig falsch angewandt und genießen doch den Nimbus der alleinseligmachenden Lösung für alle Probleme. Um LLMs aber reflektiert und erfolgreich einsetzen zu können bedarf es eines grundlegenden Verständnisses der Wirkungsweisen und der erwartbaren Ergebnisse.

Ziel des „Pen & Paper – AI Trainings“ ist es in ca. 50-60 Minuten potenziellen KI Nutzern ein Basisverständnis für die Wirkungsweise eines LLMs zu vermitteln. Dies erfolgt bewusst ohne IT-Einsatz im Methodensetting des Expertenkongresses.

Ein Expertenkongress wird in zwei Runden durchgeführt. Zuerst erarbeiten sich die Teilnehmer „Experten Kenntnisse“ zu einem Teilgebiet in themengleichen Gruppen.

In Runde zwei werden die Experten themenungleich gemischt und bearbeiten ein anspruchsvolleres Thema, das eine höhere Lernzielkategorie erfüllt und das Wissen aller Expertengruppen benötigt.

Die Ergebnisse aus Runde 2 werden im Plenum geteilt. Bearbeitungszeit für Runde 1 etwa 15-20 Minuten, für Runde 2 etwa 20-25 Minuten und etwa 15 Minuten für den Austausch nach den beiden Runden. Der gesamte Zeitbedarf liegt bei etwa 60 Minuten.

Lernziele:

Die Teilnehmenden verstehen die Arbeitsweise von Large Language Models (LLMs) im Bewerbermanagement, indem sie ihre spezifischen Rollen im Input-Management, Embedding, Intention-Management und Output-Management praxisnah anwenden. Sie analysieren Textbausteine, leiten relevante Informationen und Kompetenzen ab, transformieren diese in Embeddings, definieren Zielsetzungen und bereiten Ergebnisse adressatengerecht auf. Dabei erkennen sie die Bedeutung der Rollen im Zusammenspiel und reflektieren die Potenziale und Herausforderungen der KI-gestützten Bewerberauswahl.

Weitere Infos unter: <https://www.patternpool.de/pattern/pen-paper-ai-training/>

Rolle: Inputmanagement

Aufgabe und Bedeutung: Das Inputmanagement spielt eine zentrale Rolle im ersten Schritt der Verarbeitung von Text durch ein Large Language Model (LLM). Dessen Aufgabe ist es, den Rohtext in kleinere Einheiten, sogenannte Tokens, zu zerlegen. Diese Tokens können einzelne Wörter, Satzzeichen oder sogar Wortteile sein, je nach Komplexität des Modells. Die Zerlegung eines Textes in diese Tokens ist wichtig, weil das LLM mit dieser kleinsten Einheit arbeitet und die Bedeutung auf dieser Ebene lernt.

Im realen Kontext bedeutet das, dass das LLM ohne Inputmanagement nicht "wüsste", mit welchen Einheiten es arbeiten muss. Die korrekte Zerlegung des Textes legt den Grundstein dafür, dass das Modell später Beziehungen zwischen Wörtern herstellen kann.

Beispiel:

1. Eingabetext: "Katzenfreunde!"
2. Tokenisierung: Hier wird der Satz in Bausteine zerlegt. Meistens handelt es sich um einzelne Worte. Bei zusammengesetzten Worten wie z.B. das Wort „Katzenfreunde“ kann es in zwei Teile zerlegt werden, da das Modell möglicherweise lernt, dass „Katzen“ und „Freunde“ unterschiedliche Bedeutungen haben, die zusammen einen neuen Sinn ergeben.
3. Analyse -> Typ des Tokens als Funktionswort (geben Struktur haben wenig eigenständige Bedeutung) oder Inhaltswort (tragen den Großteil der semantischen Bedeutung eines Satzes).

Analysieren Sie den folgenden Satz und zerlegen Sie ihn in Tokens. Berücksichtigen Sie dabei die verschiedenen Wortarten und füllen Sie die untenstehende Tabelle.

Satz: „Die Katze jagt die Maus und der Hund beobachtet.“

Position	Token	Wortart	Typ
1	Die (Beispiel)	Artikel	Funktionswort

Ihr Ergebnis wird dann an das Embedding übergeben!

Rolle: Embedding

Aufgabe und Bedeutung: Embedding- Spezialistinnen und Spezialisten sind dafür verantwortlich, die Tokens (z.B. Wörter oder Wortteile) in numerische Vektoren zu übersetzen, sogenannte **Embeddings**. Diese Vektoren repräsentieren die Bedeutung der Wörter in einer mathematischen Form, sodass das LLM diese „verstehen“ und weiterverarbeiten kann. Der Embedding-Spezialist übersetzt also Sprache in Zahlen, die das Modell zur Analyse der Bedeutung verwendet.

Ohne die Embeddings könnte das Modell nicht erkennen, wie Wörter zueinander in Beziehung stehen. Diese Vektoren tragen wesentlich dazu bei, dass das Modell semantische Ähnlichkeiten oder Unterschiede zwischen Wörtern „begreift“. Ein gut trainiertes Modell wird ähnliche Bedeutungen von Wörtern durch ähnliche Zahlenmuster (Vektoren) darstellen.

Beispiel 1:

1. Token: „Katze“
2. Embedding: [0.8, 0.1, 0.3, 0.5] Das Wort „Katze“ wird in einen Vektor übersetzt. Wörter mit ähnlicher Bedeutung, wie „Hund“, hätten einen Vektor, der in einem ähnlichen Bereich liegt.

Beispiel 2:

1. Token: „laufen“
2. Embedding: [0.3, 0.6, 0.7, 0.2] Ein Verb wie „laufen“ hätte einen Vektor, der sich von Substantiven wie „Katze“ unterscheidet, aber Verben wie „rennen“ oder „gehen“ hätten ähnliche Vektoren.

Beispiel mit 3 Dimensionen:

Beispieltabelle	Tier	Größe (<i>Dimension</i>)	Lebensraum	Fortbewegungs- geschwindigkeit
	Blauwal	1.0 (<i>Wert</i>)	1.0	0.3
	Karpfen	0.2	0.0	0.2
	Blauhai	0.6	1.0	0.8
	Goldfisch	0.1	0.0	0.1

Aufgabe für die Rolle Embedding:

Erstellen Sie eine Embedding-Matrix für die folgenden Tokens, die nicht nur die Ähnlichkeit, sondern auch die semantischen Beziehungen zwischen ihnen berücksichtigt. Es sind 2 Dimensionen für die Vektoren vorgegeben. Sie können, falls aus Sicht des Teams nötig, eine dritte Dimension entwickeln und eintragen. Ordnen Sie den Token in der jeweiligen Dimension Werte von 0 bis 1 zu. Berücksichtigen Sie auch, dass ähnliche Bedeutungen ähnliche Werte in den Vektoren haben sollten.

Tokens: „Katze“, „Hund“, „Maus“

Token	Dimension und Wert		
	Größe	Geschwindigk.	
Katze			
Hund			
Maus			

Es gibt auch Verben als Tokens: „jagen“, „beobachten“

Token	Dimension und Wert		
	Aktivität	Geschwindigkeit	
jagen			
beobachten			

Ihr Ergebnis wird dann an das Attentionmanagement übergeben!

Rolle: Selbstaufmerksamkeitsmechanik (Attentionmanagement)

Aufgabe und Bedeutung: Attentionmanagerinnen und -manager überwachen die **Aufmerksamkeit** des Modells. Dies ist der Teil des Modells, der festlegt, welche Wörter in einem Satz für die Interpretation besonders wichtig sind. In der Selbstaufmerksamkeit berechnet das Modell, wie stark ein Wort auf ein anderes „aufmerksam“ sein sollte, um die Bedeutung richtig zu erfassen. Es geht also darum, festzustellen, wie Wörter miteinander in Beziehung stehen und welche Wörter besonders relevant für das Verständnis eines Satzes sind.

Das Modell könnte zum Beispiel feststellen, dass in dem Satz „Die Katze sitzt auf der Matte“ die Wörter „Katze“ und „Matte“ wichtiger sind als „der“ oder „auf“, da sie die Hauptbedeutung des Satzes tragen.

Beispiel 1:

1. Satz: „Der Hund bellt laut.“
2. Wichtige Wörter: „Hund“ und „bellt“ sind hier zentral, da sie die Hauptinformation des Satzes tragen. Das Wort „laut“ ist weniger wichtig, aber es beeinflusst die Bedeutung.

Beispiel 2:

1. Satz: „Die Sonne scheint hell.“
2. Das Modell könnte seine Aufmerksamkeit auf „Sonne“ und „scheint“ fokussieren, während „hell“ eine zusätzliche Information liefert.

Analysieren Sie den Satz „Die Katze jagt die Maus und der Hund beobachtet.“ und weisen Sie den Wörtern Aufmerksamkeitswerte zu. Berücksichtigen Sie dabei den Kontext und die Rolle jedes Wortes in Bezug auf die zentrale Handlung des Satzes. Erstelle eine Tabelle mit Token, Aufmerksamkeitswert (0 bis 1.0) und einer kurzen Begründung für den Wert.

TIPP:

Schöpfen Sie das Spektrum von 0,1 bis 1,0 aus!
Das vereinfacht Abstufungen und Unterscheidungen.

Token	Aufmerksamkeit / Attention	Begründung
Die (Beispiel)	0.1	Artikel hat geringe Relevanz für die Hauptaussage.
Katze		
jagt		
die		
Maus		
und		
der		
Hund		
beobachtet		

Ihr Ergebnis wird dann an das Vorhersagemanagement übergeben!

Rolle: Vorhersage-/Outputmanagement

Aufgabe und Bedeutung: Die Vorhersagemanagerinnen und -manager haben die Aufgabe, das nächste Wort in einem Satz oder Textabschnitt vorherzusagen. Nachdem die Tokens durch den Inputmanager zerlegt, durch den Embedding-Spezialisten in Vektoren übersetzt und durch den Attentionmanager analysiert wurden, trifft der Vorhersagemanager die Entscheidung, welches Wort wahrscheinlich als Nächstes folgt. Dies geschieht, indem er die Wahrscheinlichkeiten für möglichen nächste Wörter berechnet.

Das Vorhersage-Management ist besonders wichtig, weil es dafür sorgt, dass das Modell „flüssig“ sprechen oder Text generieren kann. Dessen Aufgabe ist es, basierend auf dem bisherigen Kontext die beste Vorhersage zu treffen.

Beispiel 1:

1. Satzbeginn: „Die Katze...“
2. Mögliche Vorhersagen: „läuft“, „schläft“, „springt“. Der Vorhersage-Manager entscheidet auf Basis der bisherigen Wörter und der Beziehung zwischen ihnen, dass „schläft“ am wahrscheinlichsten als Nächstes kommt.

Beispiel 2:

1. Satzbeginn: „Der Vogel...“

Mögliche Vorhersagen: „fliegt“, „singt“, „ruht“. Basierend auf dem Kontext entscheidet der Vorhersage-Manager, dass „fliegt“ das wahrscheinlichste nächste Wort ist.

Analysieren Sie den Satzanfang „Der Hund erhebt sich und ...“ und formulieren Sie drei plausible Vorhersagen für das nächsten Wort, wobei Sie auch jeweils die Wahrscheinlichkeit, mit der dieses Wort erwartet wird, angeben sollen (von 0 bis 1). Erklären Sie die Wahl Ihrer Vorhersagen.

Der Hund erhebt sich und...	Wahrscheinlichkeit	Erklärung

Weiterführende Aufgabe für Expertenkombi

Aufbauend auf Ihrer letzten Analyse, bei der Sie einfache Sätze mit Tieren verarbeitet haben, sollen Sie nun an einem Large Language Model zur Bewertung von Bewerbern für eine Führungsposition im Personalmanagement eines Unternehmens arbeiten. Ihr Ziel ist es, anhand von Textbausteinen aus Bewerbungsunterlagen eine KI-gestützte Analyse zu simulieren, die **relevante Skills** erkennt und die Eignung eines Kandidaten einschätzt.

Textbausteine:

1. „Frau Dr. Becker spielte mehrere Jahre als Spielführerin in der Frauenfußballmannschaft unserer Klinik.“
2. „Herr Schmidt errang mehrere bayerische Meister- und Vizemeistertitel im Schachspiel.“

Bearbeiten Sie in Ihren jeweiligen Rollen folgende Punkte und arbeiten Sie als Gruppe zusammen, um die Ergebnisse zu erzeugen:

1. Input-Management:

Zerlegen Sie die bereitgestellten Textbausteine in Tokens. Achten Sie auf eine angemessene Granularität (nicht zu klein, nicht zu grob). Ihre Aufgabe ist die formale Strukturierung der Eingabe – noch ohne Bewertung oder Gewichtung.

2. Embedding:

Wählen Sie 3–5 inhaltlich tragende Tokens aus (z. B. Substantive oder Verben - meist sind 3-5 Token für den Erkenntniseffekt dieser Übung hinreichend)

3. Selbstaufmerksamkeitsmechanik / Attentionmanagement:

Gewichten Sie die durch das Embedding repräsentierten Tokens hinsichtlich ihrer Relevanz für die konkrete Fragestellung (z. B. Eignung für eine Teamleitungsposition). Vergeben Sie Gewichtungswerte.

4. Vorhersage- / Outputmanagement

Entwerfen Sie eine einfache Darstellungsweise (z. B. eine Übersichtstabelle oder eine Stärken-/Schwächen-Analyse) und erläutern Sie, wie dies der Personalabteilung hilft.

Bearbeitungszeit ca. 20 Minuten

Bearbeitungsvorlage

Inputmanagement:

Tipp:

Denken Sie laut bei der Auswahl der relevanten Tokens – damit ermöglichen Sie Ihren Partnern das Vorausdenken!

Das Inputmanagement ist der Einstieg in die Aufgabe alle anderen Experten warten auf Ihren Input.

Ausgewählte, relevante Token:

Token

Wortart

Typ

Embedding.

Entscheiden Sie sich für besonders relevante Tokens (treffen Sie eine Auswahl meist sind 3-5 Token für den Erkenntniseffekt dieser Übung hinreichend) und weisen den vorgegebenen Dimensionen Werte zu (Ausprägung (0 bis 1.0))?

Relevante Tokens (inkl Dimension und Ausprägung (0 bis 1,0)):

		<u>Wert der Dimension</u>		
	<u>B für Dr Becker oder S für Herrn Schmidt</u>	<u>Führungspotential</u>	<u>Strategisches Denken</u>	<u>Team- orientierung</u>
<u>Token 1</u>				
<u>Token 2</u>				
<u>Token 3</u>				
<u>Token 4</u>				
<u>Token 5</u>				
<u>Token 6</u>				
<u>Token 7</u>				

Selbstaufmerksamkeitsmechanik (Attentionmanagement)

Welche Tokens aus der Liste des Embeddings sind für die Aufgabe besonders relevant. Übernehmen Sie die relevanten Tokens vom Embedding und legen Sie deren Aufmerksamkeitswert fest.

Token _____.

Aufmerksamkeitswert Frau Dr. Becker (0,0 bis 1.0): _____

Token _____.

Aufmerksamkeitswert Frau Dr. Becker (0,0 bis 1.0): _____

Token _____.

Aufmerksamkeitswert Frau Dr. Becker (0,0 bis 1.0): _____

Tokens _____.

Aufmerksamkeitswert Frau Dr. Becker (0,0 bis 1.0): _____

Token _____.

Aufmerksamkeitswert Herr Schmidt (0,0 bis 1.0): _____

Token _____.

Aufmerksamkeitswert Herr Schmidt (0,0 bis 1.0): _____

Token _____.

Aufmerksamkeitswert Herr Schmidt (0,0 bis 1.0): _____

Token _____.

Aufmerksamkeitswert Herr Schmidt (0,0 bis 1.0): _____

Vorhersage/Outputmanagement

Der Output-Manager erstellt eine Übersicht der relevanten Ergebnisse und deren Nutzen für die Personalabteilung in Form z.B. einer Tabelle basierende auf den ausgewählten Tokens und deren Aufmerksamkeitswerten des Attentionmanagements.

Beispiel:

<u>Kandidat</u>	<u>Schlüssel- info 1</u>	<u>Schlüssel- info 2</u>	<u>Schlüssel- info 3</u>	<u>Schlüssel- info 4</u>	<u>Schlüssel- info 5</u>
<u>Bevorzugtes Verhalten des Hundes, wenn eine Katze vorbei rennt,</u>	<u>Bellen -> mittel</u>	<u>Nachjagen - > hoch</u>	<u>Sitzenbleiben -> gering</u>		

<u>Kandidat</u>	<u>Schlüssel-info 1 -> Bewertung/ Nutzen</u>	<u>Schlüssel-info 2 -> Bewertung/ Nutzen</u>	<u>Schlüssel-info 3 -> Bewertung/ Nutzen</u>	<u>Schlüssel- info 4 -> Bewertung/ Nutzen</u>	<u>Schlüssel- info 5 -> Bewertung/ Nutzen</u>
<u>Frau Dr. Becker</u>					
<u>Herr Schmidt</u>					

Abschließende Bewertung:

Lösungsbeispiele:

Rolle Inputmanagement:

Token	Wortart	Anzahl der Buchstaben
Die	Artikel	3
Katze	Substantiv	5
jagt	Verb	4
die	Artikel	3
Maus	Substantiv	4
und	Konjunktion	3
der	Artikel	3
Hund	Substantiv	4
beobachtet	Verb	11
.	Satzzeichen	

Rolle Embedding:

Token	Dimension und Wert		
	Geschwindigkeit.	Größe	Karnivor
Katze	0.4	0.3	1.0
Hunde	0.4	0.5	1.0
Maus	0.3	0.1	0.0

	Aktivität	Geschwindigkeit	
laufen	0.9	0.7	
beobachten	0.2	0.1	

Rolle Selbstaufmerksamkeitsmechanik:

Token	Aufmerksamkeitswert	Begründung
Die	0.1	Artikel hat geringe Relevanz für die Hauptaussage.
Katze	0.6	Zentral für die Handlung, da sie die jagende Figur ist.
jagt	0.8	Wichtig, da es die Hauptaktion beschreibt.
die	0.1	Geringe Relevanz, wie bei „Die“.
Maus	0.7	Zentrale Figur als Ziel des Jagens.
und	0.05	Verbindet die Teilsätze, hat aber geringe Relevanz.
der	0.1	Artikel hat geringe Relevanz.
Hund	0.4	Beobachtet und ist relevant, aber nicht aktiv.
beobachtet	0.5	Wichtig für den Kontext, trägt zur Handlung bei.

Rolle Vorhersage Manager:

- „bellt“ (Wahrscheinlichkeit: 0.6)
Erklärung: Hunde bellen oft wenn sie sich herausgefordert sehen.
- „rennt los“ (Wahrscheinlichkeit: 0.4)
Erklärung: Eine plausible Möglichkeit, da Hunde einen ausgeprägten Jagdinstinkt haben.
- „legt sich“ (Wahrscheinlichkeit: 0.1)
Erklärung: Weniger wahrscheinlich, da dies nicht die häufigste Position ist, aber immer noch eine Möglichkeit.

Lösungsbeispiel weiterführende Expertenaufgabe:

Rahmen:

Bewerbersauswahl für eine Teamleitungsposition im Krankenhaus.

Textbausteine:

1. „Frau Dr. Becker spielte mehrere Jahre als Spielführerin in der Frauenfußballmannschaft unserer Klinik.“
2. „Herr Schmidt errang mehrere bayerische Meister- und Vizemeistertitel im Schachspiel.“

Input-Management

(Formale Struktur – keine Bewertung, keine Relevanzentscheidung)

Tokenisierung – Text 1

Frau | Dr. | Becker | spielte | mehrere | Jahre | als | Spielführerin | in | der |
Frauenfußballmannschaft | unserer | Klinik | .

Tokenisierung – Text 2

Herr | Schmidt | errang | mehrere | bayerische | Meister- | und | Vizemeistertitel | im | Schachspiel
| .

Strukturelle Kennzeichnung (Beispielauszug)

Token	Wortart	Typ
Spielführerin	Substantiv	Inhaltswort
Frauenfußballmannschaft	Substantiv	Inhaltswort
Meistertitel	Substantiv	Inhaltswort
Schachspiel	Substantiv	Inhaltswort
mehrere	Quantor	Funktionswort

Embedding

(Bedeutungsrepräsentation – keine Gewichtung)

Vorgegebene Dimensionen:

1. Führungspotenzial
2. Strategisches Denken
3. Teamorientierung

Embedding-Matrix (Beispielwerte)

Token	Führung Strategie Team		
Spielführerin	0.9	0.4	0.8
Frauenfußballmannschaft	0.4	0.2	0.9
Klinik	0.5	0.3	0.7
Meistertitel	0.6	0.7	0.3
Vizemeistertitel	0.5	0.7	0.3
Schachspiel	0.2	0.9	0.1

Funktionswörter → 0 in allen Dimensionen.

Erklärung:

- Spielführerin → stark Führung + Team
- Schachspiel → stark Strategie
- Meistertitel → Leistung + Wettbewerb

Keine Gewichtung im Kontext der Stelle – nur Bedeutungsdarstellung.

Attention-Management

(Kontextabhängige Gewichtung – keine Zieldefinition)

Kontext:

Eignung für Teamleitungsposition im Krankenhaus.

Attention-Gewichte – Kandidatin Becker

Token	Attention
Spielführerin	0.9
Frauenfußballmannschaft	0.7
Klinik	0.6
spielte	0.4
mehrere Jahre	0.5

Attention-Gewichte – Kandidat Schmidt

Token	Attention
Schachspiel	0.8
Meistertitel	0.7
Vizemeistertitel	0.6
errang	0.4

Begründung:

- Für Teamleitung sind Führung und Teambezug besonders relevant.
- Strategisches Denken ist relevant, aber sekundär.
- Wettbewerbsleistung ist unterstützend, aber nicht zentral.

Attention priorisiert Informationen – sie entscheidet noch nicht.

Vorhersage / Output-Management

(Aggregierte Bewertung und Empfehlung)

Aus gewichteten Embeddings ergibt sich:

Bewertungsübersicht

Kandidat	Führung	Strategie	Team	Gesamtbewertung
Becker	Hoch	Mittel	Hoch	Sehr geeignet
Schmidt	Niedrig	Hoch	Niedrig	Fachlich geeignet

Beispielhafter KI-Output

„Auf Basis der gewichteten Analyse weist Frau Dr. Becker ein ausgeprägtes Führungspotenzial mit starkem Teambezug auf. Herr Schmidt zeigt hohe strategische Kompetenz, jedoch weniger explizite Hinweise auf Führungserfahrung im Teamkontext.“

Zusammenfassung

Die Aufgabe fordert die Gruppenmitglieder, ihre spezifischen Rollen zu nutzen, um die bereitgestellten Textbausteine zu analysieren, in Embeddings zu übersetzen, Intentionen zu definieren und Ergebnisse aufzubereiten. So wird das Zusammenspiel der Rollen in einem LLM praxisnah vermittelt.