



Ralf Walther
Geschäftsführer & Gründer
mindUp Web + Intelligence GmbH

mindUp
Web + Intelligence GmbH

Vorstand & Gründungsmitglied
cyberLAGO e.V.

cyberLAGO
digital competence network

KI-LAB Bodensee
KI-Kompetenz für die Region

KI-Kompetenz aufbauen
KI-Anwendungen erproben
KI-Experten finden

KI / LLMs im
Wissensmanagement



Das KI-Universum
wächst weiter



Stärken
und
Schwächen



Wohin geht
die Reise?



Next Level
LLMs

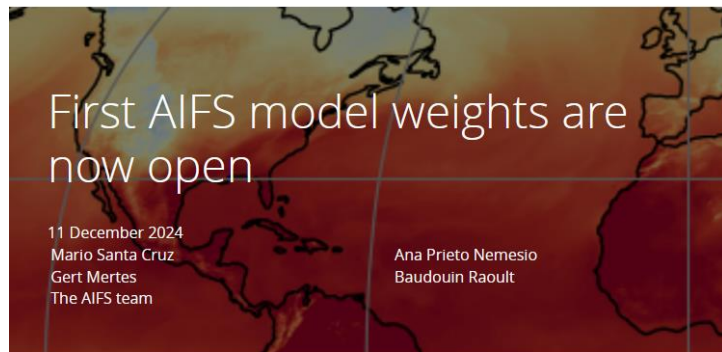




European Centre for Medium-Range Weather Forecasts



Who are we | What we do | Jobs | **Media centre** | Suppliers | Location



◀ [View all AIFS blog posts](#)

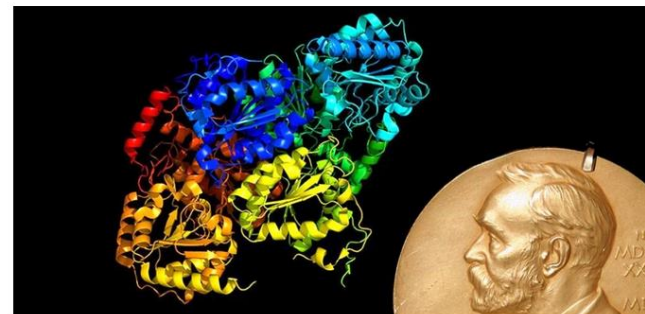
We are excited to announce a key step in making our data-driven weather model even more accessible to a broader audience of developers, scientists, and weather enthusiasts. Already you can access real-time [charts](#) and [data](#) for the AIFS, and now the latest deterministic version of the AIFS (AIFS-single), v0.2.1, is available [on Hugging Face](#) so you can run the AIFS on your own device. Hugging Face is a widely used platform and community hub for sharing machine learning models. It provides an open ecosystem where developers can share, access, and experiment with state-of-the-art models across disciplines. By hosting

Chemie

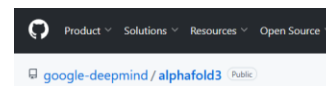
Chemie-Nobelpreis für KI-gestützte Proteinforschung

Entwickler der KI AlphaFold und Pionier künstlicher Proteine ausgezeichnet

9. Oktober 2024, Lesezeit: 4 Min.



Den Chemie-Nobelpreis 2024 erhalten drei Pioniere der computergestützten Proteinforschung. © gemeinfrei



AlphaFold 3

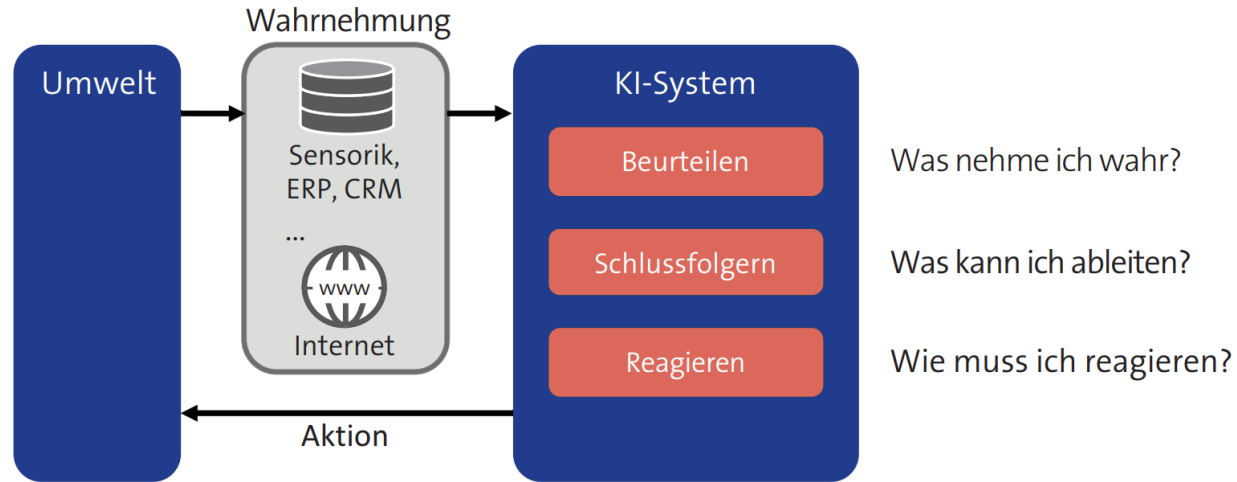
This package provides an implementation of the inference pipeline of AlphaFold 3. See below for how to access the model parameters. You may only use AlphaFold 3 model parameters if received directly from Google. Use is subject to these [terms of use](#).

Any publication that discloses findings arising from using this source code, the model parameters or outputs produced by those should [cite](#) the [Accurate structure prediction of biomolecular interactions with AlphaFold 3](#) paper.

Please also refer to the Supplementary Information for a detailed description of the method.

AlphaFold 3 is also available at [alphafoldserver.com](#) for non-commercial use, though with a more limited set of ligands and covalent modifications.





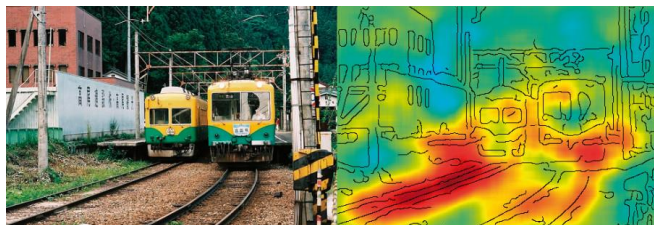


Was soll mit dem Einsatz von KI erreicht werden?

- Effizienzsteigerung
 - Automatisierung von wiederkehrenden Aufgaben
 - Analyse von großen Datenmengen
 - Konzentration der Akteure auf komplexe Aufgaben, die (noch) nicht von KI gelöst werden können
- Qualitätsverbesserung
 - Validierung/Standardisierung von Ergebnissen
 - Messbarkeit
- Kommunikation/Kollaboration



- Überhöhte Erwartungen, unterschätzte Unschärfe
- Falsche Zieldefinition aufgrund Unkenntnis der Möglichkeiten
- Zu wenig oder ungeeignete Daten: Ohne Daten kann keine KI-basierte Lösung erfolgreich sein
- Was nicht in den Daten steckt, kann nicht von der KI gelernt werden
- Interpretierbarkeit des Modells/der Ergebnisse



Quelle: <https://www.hhi.fraunhofer.de/presse-medien/nachrichten/2019/wie-intelligent-ist-kuenstliche-intelligenz.html>



Realistische Darstellung von mehreren Armbanduhren auf einem modernen Tisch, die alle die Zeit 17:35 zeigen.

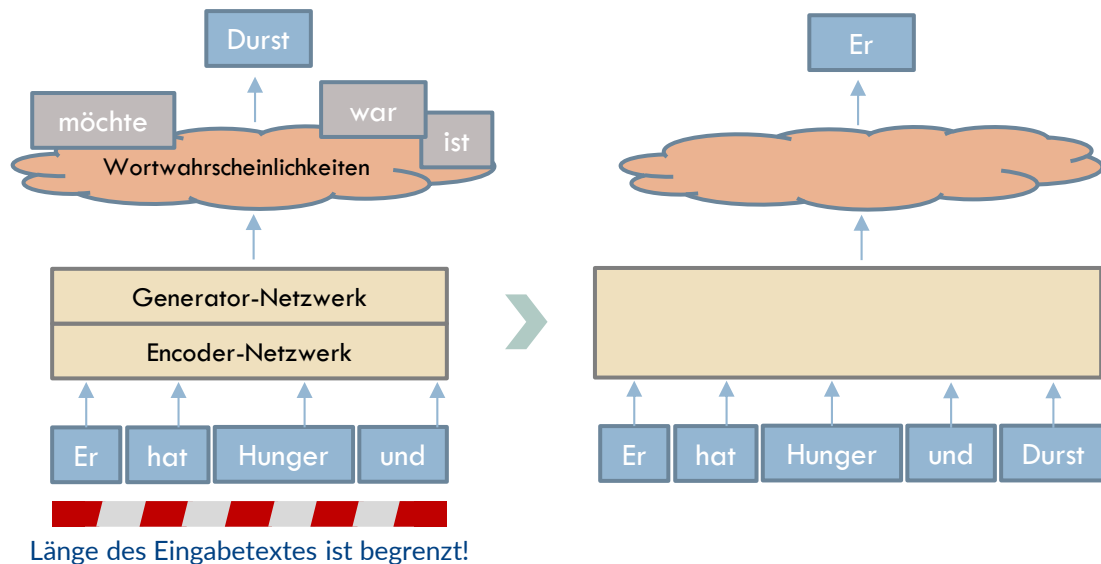


Realistische Darstellung eines Fortbildungskurses für Zeichnen. An einem Tisch sitzen mehrere Linkshänder, die mit der linken Hand eine Zeichnung anfertigen.

Finde den Fehler!



Model	Lab	Playground	Parameters (B)	Tokens trained (B)	Ratio Tokens/Params (Chinchilla scaling 20:1)	ALScore ("ALScore" is Sqr Root of	MMLU	MMLU -Pro	GPQA	Training dataset	Announced ▼	Public?	Paper / Arch Repo	Tags	Notes	
DeepSeek-V3 0324	DeepSeek-AI	https://chat.deepseek.com	685	14800	22:1	10.6	81.2	68.4	synthetic, web-scale	Mar/2025	🟢	https://huggingface.co/deepseek-ai/DeepSeek-V3	MoE	SOTA	Non-reasoning. Significant increase in benchmark performance compared to original V3	
Llama-3.3-Nemotron	NVIDIA	https://huggingface.co/nvidia/Llama-3.3-Nemotron	49	15040	307:1	2.9		66.67	web-scale	Mar/2025	🟢	https://huggingface.co/nvidia/Llama-3.3-Nemotron	Dense	Reasoning	Meta Llama-3.3-70B-Instruct derivative "that is post trained for reasoning, human chat, & reasoning"	
EXAONE Deep	LG	https://huggingface.co/exaone/exaone-deep	32	6500	204:1	1.5	83	74	66.1	web-scale	Mar/2025	🟢	https://arxiv.org/abs/2503.10940	Dense	Reasoning	"EXAONE"="Expert AI for EveryONE". Training tokens/ratio dropped from EXAONE-3 7.8
Mistral Small 3.1	Mistral	https://huggingface.co/mistralai/Mistral-Small-3.1	24	8000	334:1	1.5	81.01	56.03	37.5	web-scale	Mar/2025	🟢	https://mistral.ai/news/mistral-small-3.1/	Dense		"Mistral Small 3.1 (2503) adds state-of-the-art vision understanding and enhances long context"
ERNIE 4.5	Baidu	https://qianwen.lingji.com.cn/qianwen45						79	55	synthetic, web-scale	Mar/2025	🟢	https://www.baidubce.com/en/Products/ERNIE-4.5	MoE		Announce: https://x.com/Baidu_Inc/status/1901094083508220035
X1	Baidu	https://qianwen.lingji.com.cn/qianwenx1								synthetic, web-scale	Mar/2025	🟢	https://www.baidubce.com/en/Products/X1	MoE	Reasoning	
OLMo 2 32B	Allen AI	https://playground.allenai.org/	32	6400	200:1	1.5	78			synthetic, web-scale	Mar/2025	🟢	https://allenai.org/ollmo	Dense		"the first fully-open model (all data, code, weights, and details are freely available) to our knowledge"
Command A	Cohere	https://dashboards.cohere.com/	111	8000	73:1	3.1	85			synthetic, web-scale	Mar/2025	🟢	https://huggingface.co/cohere/Command-A	Dense		Context=256k. "Command A is an open weights research release of a 111 billion parameter model"
Gemini Robotics	Google DeepMind		200	20000	100:1	6.7		79.1	64.7	synthetic, web-scale	Mar/2025	🔴	https://storage.googleapis.com/generativeai-downloads/public/gemini_robotics_v1.0.0_20250307.pdf	MoE		Gemini 2.0 Pro (cloud). "The second model is Gemini Robotics, a state-of-the-art Vision-Language model"
Gemini Robotics-ER	Google DeepMind		30	30000	1,000:1	3.2	87	76.4	62.1	synthetic, web-scale	Mar/2025	🔴	https://storage.googleapis.com/generativeai-downloads/public/gemini_robotics_er_v1.0.0_20250307.pdf	MoE		Gemini 2.0 Flash (on device). "The first model is Gemini Robotics-ER, a VLM with strong reasoning capabilities"
Gemma 3	Google DeepMind	https://huggingface.co/google/gemma-3	27	14000	519:1	2.0	78.6	67.5	42.4	web-scale	Mar/2025	🟢	https://huggingface.co/google/gemma-3	Dense		Trained on 1T more tokens than Gemma 2. "Introduces vision understanding abilities, a reasoning model, and a MoE model"
Reka Flash 3	Reka AI	https://huggingface.co/reka/reka-flash-3	21	5000	239:1	1.1		65	61.1	web-scale	Mar/2025	🟢	https://www.reka.ai/	Dense	Reasoning	"performs competitively with proprietary models such as OpenAI o1-mini, making it a great choice for a wide range of applications"
QwQ-32B	Alibaba	https://huggingface.co/Qwen/QwQ-32B	32	18000	563:1	2.5				synthetic, web-scale	Mar/2025	🟢	https://qwenlm.github.io/	Dense	Reasoning	Update to QwQ-32B-Preview released Nov/2024. Scores 1/5 on latest ALPROMPT 2024 H
Jamba 1.6	AI21	https://huggingface.co/aion/jamba-1.6	398	8000	21:1	5.9	81.2	53.5	36.9	web-scale	Mar/2025	🟢	https://ai21.com/jamba	MoE		"The AI21 Jamba 1.6 family of models is state-of-the-art, hybrid SSM-Transformer instruction-tuned LLMs"
Instella-3B	AMD	https://huggingface.co/instella/instella-3b	3	4160	1,387:1	0.4	58.31		30.13	web-scale	Mar/2025	🟢	https://rocm.ai/instella	Dense		"trained from scratch on AMD Instinct™ MI300X GPUs. Instella models outperform existing models on a wide range of benchmarks"
Babel-83B	Alibaba	https://huggingface.co/Babel/Babel-83B	83	15000	181:1	3.7				synthetic, web-scale	Mar/2025	🟢	https://arxiv.org/abs/2503.10940	Dense		"top 25 languages by number of speakers, including English, Chinese, Hindi, Spanish, Arabic, and more"
Granite-3.2-8B-Instr	IBM	https://huggingface.co/ibm/granite-3.2-8b-instruct	8	12000	1,500:1	1.0	66.79			web-scale	Feb/2025	🟢	https://www.ibm.com/press/ibm-granite-3.2	Dense	Reasoning	"The new Granite 3.2 8B Instruct [offers] experimental chain-of-thought reasoning capabilities"
C4AI Command R7B	Cohere	https://huggingface.co/cohere/command-r7b	7	2000	286:1	0.4		29.4	7.9	web-scale	Feb/2025	🟢	https://huggingface.co/cohere/command-r7b	Dense		"C4AI Command R7B Arabic is an open weights research release of a 7 billion parameter model"
GPT-4.5	OpenAI	https://chat.openai.com/	5400	114000	22:1	82.7	89.6		71.4	synthetic, web-scale	Feb/2025	🟢	https://openai.com/index/introducing-gpt-4.5/	MoE	SOTA	"Our largest and best model for chat" https://openai.com/index/introducing-gpt-4.5/ / "C4AI Command R7B"
Hunyuan T1	Tencent	https://cloud.tencent.com/ai/hunyuan-t1	389	7000	18:1	5.5		87.2	69.3	synthetic, web-scale	Feb/2025	🟢	https://img.tencent.com/ai/hunyuan-t1	MoE	Reasoning	"Based on Turbo S, by introducing technologies such as long thinking chains, retrieval augmentation, and more"
Hunyuan Turbo S	Tencent	https://cloud.tencent.com/ai/hunyuan-turbo-s	389	7000	18:1	5.5	89.5	79	57.5	synthetic, web-scale	Feb/2025	🟢	https://img.tencent.com/ai/hunyuan-turbo-s	MoE		Fast thinking ("Instant reply"). "This is also the first time in the industry that the Mamba architecture is used"
Phi-4-multimodal	Microsoft	https://huggingface.co/microsoft/phi-4-multimodal	5.6	6100	1,090:1	0.6				synthetic, web-scale	Feb/2025	🟢	https://huggingface.co/microsoft/phi-4-multimodal	Dense		"Training data: 5T tokens, 2.3M speech hours, and 1.1T image-text tokens" Announce: https://microsoft.com/en-us/ai/phi-4
Phi-4-mini	Microsoft	https://huggingface.co/microsoft/phi-4-mini	3.8	5000	1,316:1	0.5	67.3	52.8	30.4	synthetic, web-scale	Feb/2025	🟢	https://huggingface.co/microsoft/phi-4-mini	Dense		"Phi-4-mini's training data includes a wide variety of sources, totaling 5 trillion tokens, a mix of synthetic and real data"
Mercury Coder Sma	Inception	https://chat.inception.ai/	40	5000	125:1	1.5				synthetic, web-scale	Feb/2025	🟢	https://www.inception.ai/	Dense		Diffusion large language model (dLLM). Very low 'IQ' performance (0/5 on all ALPROMPT: 2024 H)
QwQ-Max-Preview	Alibaba	https://chat.qwenlm.com/	325	20000	62:1	8.5				synthetic, web-scale	Feb/2025	🟢	https://qwenlm.github.io/	Dense	Reasoning	"As a sneak peek into our upcoming QwQ-Max release, this version offers a glimpse of its capabilities"
Claude 3.7 Sonnet	Anthropic	https://claude.ai/	175	20000	115:1	6.2			84.8	synthetic, web-scale	Feb/2025	🟢	https://anthropic.com/news/claude-3.7-sonnet	Dense	Reasoning SOTA	Knowledge cutoff now November 2024 (was April 2024). "the first hybrid reasoning model"
Moonlight	Moonshot AI	https://huggingface.co/moonshotai/moonlight	16	5700	357:1	1.0	70	42.4		synthetic, web-scale	Feb/2025	🟢	https://github.com/moonshotai/moonlight	MoE		"Scaling law experiments indicate that Muon achieves ~ 2x computational efficiency compared to Llama"
S2	Figure		7	2000	286:1	0.4				synthetic, web-scale	Feb/2025	🔴	https://www.figure.ai/	Dense		Likely based on OpenVLA 7B (Jun/2024, based on Llama 2 7B) or Molmo 7B-O (Sep/2024)
S1	Figure		0.08	1	13:1	0.0				special	Feb/2025	🔴	https://www.figure.ai/	Dense		"high quality, multi-robot, multi-operator dataset of diverse teleoperated behaviors, ~50C"
Baichuan-M1-14B	Baichuan	https://huggingface.co/baichuan/baichuan-m1-14b	14	20000	1,429:1	1.8				synthetic, web-scale	Feb/2025	🟢	https://arxiv.org/abs/2503.10940	Dense		Medical LLM. Huge increase to 20T tokens for 14B params standard.
Evo 2	Arc Institute	https://github.com/arc-institute/evo-2	40	8800	220:1	2.0				special	Feb/2025	🟢	https://github.com/arc-institute/evo-2	Dense		"Evo 2 is a state of the art DNA language model for long context modeling and design. Ev"
R1 1776	Perplexity	https://huggingface.co/perplexity/r1-1776	685	14800	22:1	10.6	90.5			synthetic, web-scale	Feb/2025	🟢	https://www.perplexity.ai/	MoE	Reasoning	Censorship reduced, based on DeepSeek-R1.
Grok-3	xAI	https://grok.com/	928	36200	40:1	19.3		79.9	84.6	synthetic, web-scale	Feb/2025	🟢	https://x.ai/blog/grok-3	MoE	Reasoning SOTA	https://x.ai/blog/grok-3 My full analysis: https://lifearchitect.ai/whats-in-grok/
Mistral Saba	Mistral	https://console.mistral.ai/	24	8000	334:1	1.5	81			web-scale	Feb/2025	🟢	https://mistral.ai/news/mistral-saba/	Dense		"Mistral Saba is a 24B parameter model trained on meticulously curated datasets from a variety of sources"
Salamandra	Barcelona Superco	https://github.com/bsc-sal/salamandra	40	9000	225:1	2.0				web-scale	Feb/2025	🟢	https://arxiv.org/abs/2503.10940	Dense		"The final [pre-training] dataset is composed of 55.51% FineWeb-Edu, 25.32% Colossal (C"



Durch GPT findet eine probabilistische, kontextbezogene **Wortvorhersage** statt.



[Overview](#) [Documentation](#) [Examples](#) [Playground](#)

Playground

Er hat Hunger und **Durst**

Er sollte etwas zu essen und **trinken**. Einige gute Optionen wären:

- Frisches Obst oder Gemüse
- Eine **S**chüssel Reis oder Nudeln
- Ein S
- Ein G Hand = 21.29%
- Ein S Mah = 19.16%

g = 10.92%
le = 6.56%
Port = 6.09%
Sch = 4.94%

Total: -3.01 logprob on 1 tokens
(68.96% probability covered in top 6 logits)

Maßnahmen zur Anpassung

Es existieren verschiedene Techniken, um die Funktionalitäten des LLMs für die eigene Anwendung anzufertigen.

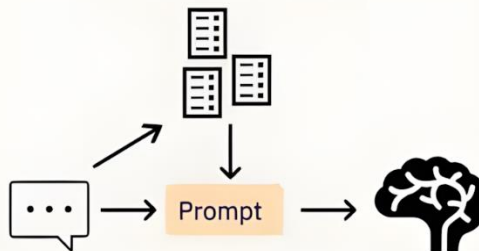
Prompt Engineering

Prompt

- Standardanweisung
- Spezielle Anweisungen
- Beispiele:
 - Beispiel 1 – Antwort 1
 - Beispiel 2 – Antwort 2
 - ...

Füge spezielle Anweisungen oder Beispiele dem Prompt hinzu

Retrieval Augmented Generation (RAG)



Stimme Input mit eigenen Daten ab und nutze passende Daten als neuen Promptkontext

Fine-Tuning



Nutze eigene Daten um ausgewählte Parameter im LLM neu zu trainieren

Grad der Komplexität / Rechnerischer Aufwand

Standard Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

Model Output

A: The answer is 27. ❌

Chain-of-Thought Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: Roger started with 5 balls. 2 cans of 3 tennis balls each is 6 tennis balls. $5 + 6 = 11$. The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

Model Output

A: The cafeteria had 23 apples originally. They used 20 to make lunch. So they had $23 - 20 = 3$. They bought 6 more apples, so they have $3 + 6 = 9$. The answer is 9. ✅

Du bist Mirko Schlaumeier, Kundenberater für den Bereich Blankstahl, Qualitätsstahl und hast eine Email-Anfrage erhalten.

Start der Emailanfrage des potenziellen Kunden:

Guten Tag Herr Schlaumeier,

wir hatten schon über meine Bestellung **telefonisch** gesprochen.

Bitte machen Sie mir doch eine **Angebot** fertig mit den besprochenen Positionen.

Vielen Dank

Hans Mustermann

Ende der Email

Eine **Suche** in der Kommunikation mit dem Kunden hat folgende Inhalte gefunden:

- Anruf 15.02.2025 14:24 Uhr von Herrn Mustermann
- er war sehr zufrieden und hat sich über die Sonderangebote bzgl. Kurzlängen informiert
- es kommen 500 Stück zu 3m Länge in Frage
- Art. Nr. 705171

RAG

Ende der Kommunikationsdaten

Eine **Suche** in der Artikeldatenbank hat folgendes Ergebnis zu den Artikelnummern ergeben:

Artikelnummer **705171**

Bezeichnung: 25CrMo4+QT+SH h9 beids. gefast

Preis: 1100 € / Tonne

Aktueller Lagerbestand: 574

Form: Rd 12

RAG

Ende des Suchergebnisses zu Artikelstammdaten

Prüfe, ob die geforderte Menge an Produkten verfügbar ist.

Sofern die geforderte Menge verfügbar ist, schreibe einen kurzen **Angebotstext**, der das Produkt beschreibt. Nutze eine geeignete Anrede und Schlußformel.

Assistant

Guten Tag Herr Mustermann,

vielen Dank für Ihren Anruf und Ihr Interesse an unseren Blankstahl-Produkten. Ich freue mich, Ihnen das folgende Angebot unterbreiten zu können.

Sie interessieren sich für 500 Stück zu je 3 Metern Länge des Artikels 705171, Bezeichnung: 25CrMo4+QT+SH h9 beidseitig gefast. Dieses Produkt bieten wir aktuell zum Preis von 1100 € pro Tonne an. Der aktuelle Lagerbestand beträgt 574 Einheiten, sodass die von Ihnen gewünschte Menge von 500 Stück erfüllt werden kann.

Sollten Sie weitere Fragen haben oder eine Bestellung aufgeben wollen, stehe ich Ihnen selbstverständlich jederzeit zur Verfügung.

Ich freue mich darauf, bald wieder von Ihnen zu hören.

Mit freundlichen Grüßen,

Mirko Schlaumeier

Kundenberater Blank- und Qualitätsstahl



how many Rs are in strawberry

OI deepseek-r1:8b

<think>

Okay, so I'm trying to figure out how many times the letter "R" appears in the word "strawberry" by breaking it down and examining each letter one by one.

The word is S-T-R-A-W-B-E-R-R-Y. Hmm, I should probably write it out in a more linear way by breaking it down:

1. S
2. T
3. R
4. A
5. W
6. B
7. E
8. R
9. R
10. Y

Wait, let me make sure I have the letters in order. Sometimes when I'm focusing on parts of a letter. Let me write it again: S, T, R, A, W, B, E, R, R, Y.

I think I'm confident with this count now. It's always good to double-check because some words are spelled differently or if I'm tired. But in this case, I'm pretty sure it's three R's.
</think>

The letter "R" appears three times in the word "strawberry."

Step-by-Step Explanation:

1. **Write Down the Word:** Start by writing out the word to visualize each letter clearly.
 - S T R A W B E R R Y
2. **Identify Each Letter Position:**
 - 1: S
 - 2: T
 - 3: R
 - 4: A
 - 5: W
 - 6: B
 - 7: E
 - 8: R
 - 9: R
 - 10: Y
3. **Count the "R"s:**
 - Position 3: R (1st occurrence)
 - Position 8: R (2nd occurrence)
 - Position 9: R (3rd occurrence)
4. **Total Count:** There are three "R" letters in the word.

Answer: The letter "R" appears three times in "strawberry."

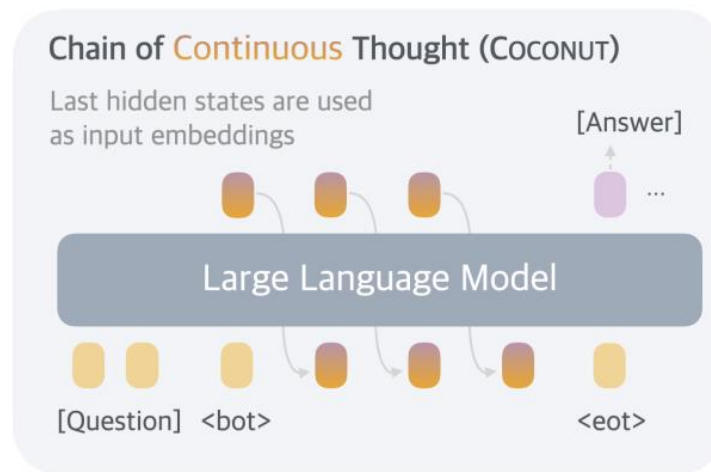
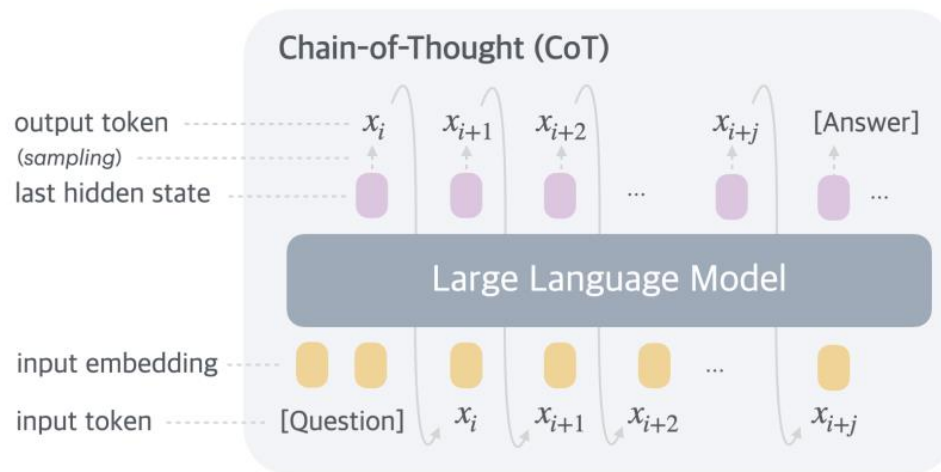
Training Large Language Models to Reason in a Continuous Latent Space

Shibo Hao^{1,2,*}, Sainbayar Sukhbaatar¹, DiJia Su¹, Xian Li¹, Zhiting Hu², Jason Weston¹, Yuandong Tian¹

¹FAIR at Meta, ²UC San Diego

*Work done at Meta

Large language models (LLMs) are restricted to reason in the "language space", where they typically express the reasoning process with a chain-of-thought (CoT) to solve a complex reasoning problem.



Chain-of-Thoughts auf Basis der **internen** Repräsentation des Ergebnisses.
Es werden **keine Wörter** für die Zwischenergebnisse generiert.

paipz.de

paipz - KI-gestützte Dokumentenverarbeitung

product name:
LD COVID-19 IgG Rapid Test 100.0% Accept Reject Edit

catalog nr:
COVID-19_G_10_EN 100.0% Accept Reject Edit

manufacturer address:
**Salzdahlumer Str 196, Geb. 1A
D-38126 Braunschweig** 100.0% Accept Reject Edit

manufacturer phone:
+49(0)531 6180654 100.0% Accept Reject Edit

manufacturer emergency_phone:
+49(0)551 19240 75.4% Accept Reject Edit

manufacturer_email:
info@lionex.de 88.6% Accept Reject Edit

product_identification:
IG COVID-19 IgG Rapid Test 100.0% Accept Reject Edit

product_identification_Art-Nr:
COVID-19_G_10_EN 96.4% Accept Reject Edit

MSDS (Material Safety Data Sheet / Sicherheitsdatenblatt) DO 619

DO 619	Gültig ab: 03.07.2017	GMH - Abschnitt: 8.5	Seiten: 1 von 10
Revision / Fassung Nr.: 3.6	Gültig ab: 28.05.2020	Produktname: LD COVID-19 IgG Rapid Test	Katalog-Nr.: COVID-19_G_10_EN

EG Material Safety Data Sheet according to
Safety Data Sheet according to Regulation (EG) 2015/830 Regulation, (EU) No. 1272/2008 (+ Subsequent ATPs) and
REACH Regulation 1907/2006 EC (+ Subsequent Regulations)

Date: 28.05.2020 Rev. 1.0

SECTION 1. Identification of the substance/mixture and of the company/undertaking

1.1. 1.1. Product identifier: LD COVID-19 IgG Rapid Test (Art.-No. COVID-19_G_10_EN)

1.2. Relevant identified uses of the substance or mixture and uses advised against

The LD COVID-19 IgG Rapid Test is a visual test for the qualitative detection of IgG antibodies to SARS-CoV-2 in human serum, plasma, or whole blood within 15 min (latest 20 min). The test is for professional in vitro diagnostic use only and intended to aid in the diagnosis of SARS-CoV-2 infection and to complement direct pathogen detection. Serology can also be used to collect epidemiological data. The product is intended for use as an IVD but can also be used for research purposes. Contains 2 components (test cassette and diluent buffer) as liquids or solid phase.

1.3. Details of the supplier of the safety data sheet

Lionex GmbH
Salzdahlumer Str. 196, Geb. 1A
D-38126 Braunschweig
Tel. +49(0)531 / 2601266
Contact person: Prof. Dr. Singh:
www.lionex.de
FAX +49(0)531 / 6180654
e-mail: info@lionex.de
Tel. +49(0)175 / 594 2291

1.4. Emergency telephone number

Germany:
Giftnformationszentrum-Nord der Länder Bremen,
Hamburg, Niedersachsen und Schleswig-Holstein
Robert-Koch-Straße 40
37075 Göttingen
Tel.: +49(0)551 / 19240

International:

Belgien / Belgium:	+32(0) 245 246	Polen / Poland:	+48 (62) 657 99 00
Bulgarien / Bulgaria:	+359 (0) 515 32 34	Portugal / Portugal:	+351 (21) 795 51 43
Dänemark / Denmark:	+45 (35) 316 060	Russische - Föderation / Russia:	+7 (95) 928 16 47
Finnland / Finland:	+358 (9) 421 977	Schweden / Sweden:	+46 (8) 736 03 84
Frankreich / France:	+33 (3) 883 737 97	Schweiz / Switzerland:	+41 (21) 251 51 51

weyp © 2025

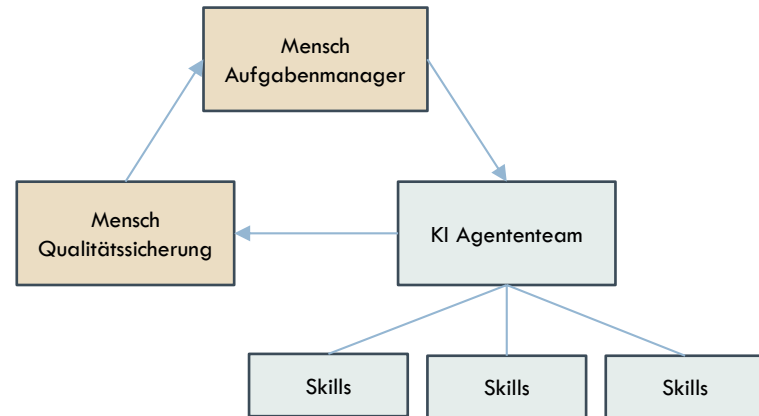
Vorteile KI-basierter Extraktion:

- Training auf bestehenden PDF-Sammlungen
- Ortsinvariant, d.h. kein Oldschool-Templating
- Berücksichtigung der Fachterminologie
- (fast) sprachunabhängig
- Automatisierungsgrad individuell einstellbar



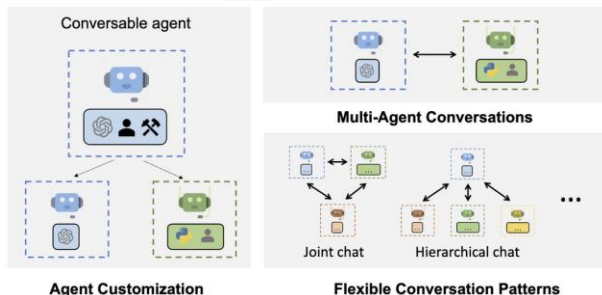
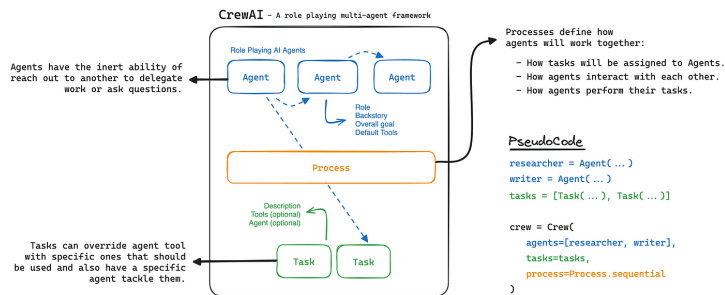
Arbeitsteilung KI

- Multi-Agenten-Systeme
- Arbeitsteilung/Kollaboration
- Verwendung von unterschiedlichen LLMs/SLMs
- Open-Source, Fully-Local
- Cloud-basierte Skalierung von Services





Home Enterprise Open Source Ecosystem Use Cases Templates Blog



Quelle: techcommunity.microsoft.com/t5/ai-azure-ai-services-blog/building-a-multimodal-multi-agent-framework-with-azure-openai/ba-p/4084007

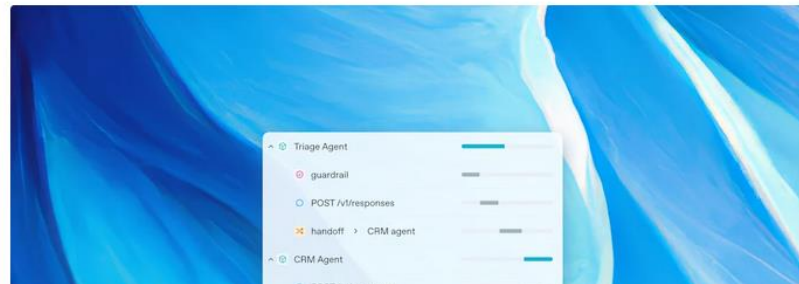
OpenAI

March 11, 2025 Product

New tools for building agents

We're evolving our platform to help developers and enterprises build useful and reliable agents.

Try in Playground



OpenAI Agents SDK

OpenAI Agents SDK

Intro
Quickstart
Examples
Documentation
Agents
Running agents
Results
Streaming
Tools

OpenAI Agents SDK

The **OpenAI Agents SDK** enables you to build agentic AI apps in a lightweight, easy-to-use package with very few abstractions. It's a production-ready upgrade of our previous experimentation for agents, **Swarm**. The Agents SDK has a very small set of primitives:

- **Agents**, which are LLMs equipped with instructions and tools
- **Handoffs**, which allow agents to delegate to other agents for specific tasks
- **Guardrails**, which enable the inputs to agents to be validated

SDK ist Open-Source!