

Large Language Models

Simon Ellershaw

PhD Candidate Institute of Health Informatics UCL

simon.ellershaw.20@ucl.ac.uk

Buzzwords

- LLM
- GPT
- Transformer
- Next token prediction
- Few shot prompting
- Chain of Thought
- Self-Consistency
- Finetuning
- Multi Modal
- Foundation Models
- RLHF
- Scaling Laws
- Open Source
- RAG
- Semantic Search
- Embeddings
- Agents
- Reasoning Model

Content

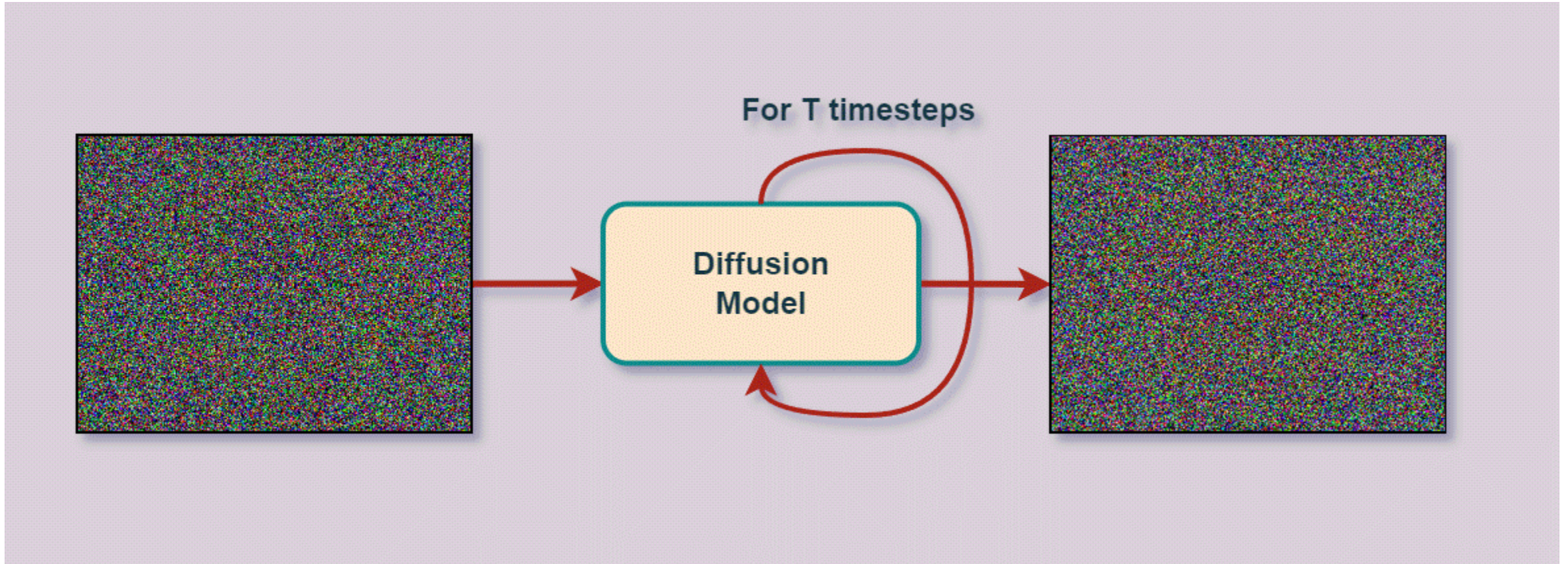
1. What is an LLM?
2. How to use an LLM
3. What does this mean for medicine

Content

1. What is an LLM?
2. How to use an LLM
3. What does this mean for medicine



Content

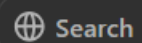


ChatGPT ▾

SI

What can I help with?

Write the answer to Prof Taylor's latest assignment



Search



Reason



Create image



Code



Summarize text



Help me write

More

So what is GPT-4?

GPT = Generative Pre-trained Transformer

“GPT-4 is a **Transformer-style model** [39] pre-trained to **predict the next token** in a document, using both publicly available data (such as internet data) and data licensed from third-party providers. The model was then fine-tuned using **Reinforcement Learning from Human Feedback** (RLHF) [40]. Given both the competitive landscape and the safety implications of large-scale models like GPT-4, this report contains **no further details** about the architecture (including model size), hardware, training compute, dataset construction, training method, or similar.”

Transformer

[CLS]

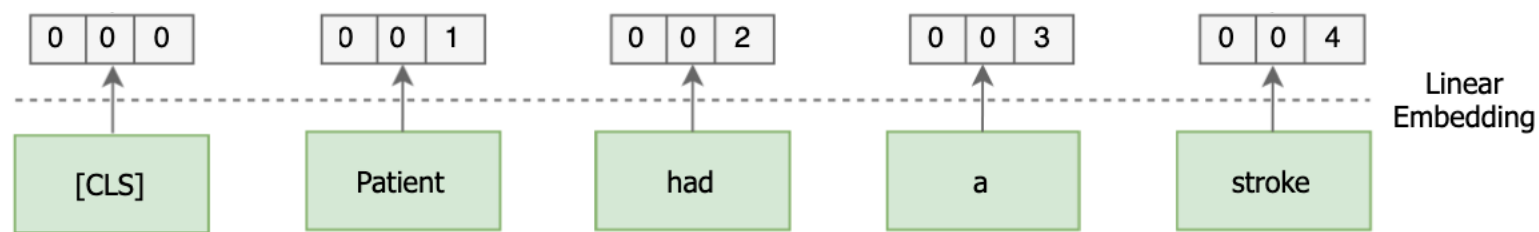
Patient

had

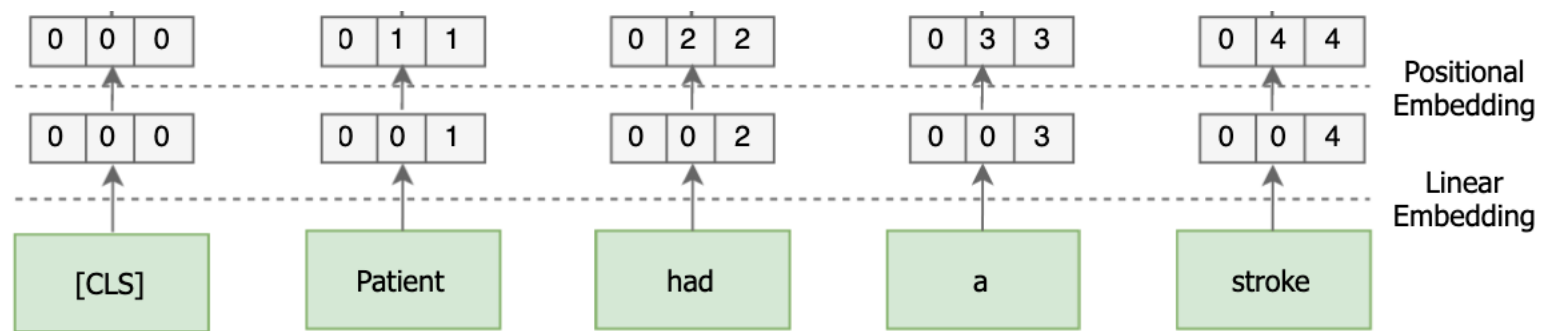
a

stroke

Transformer

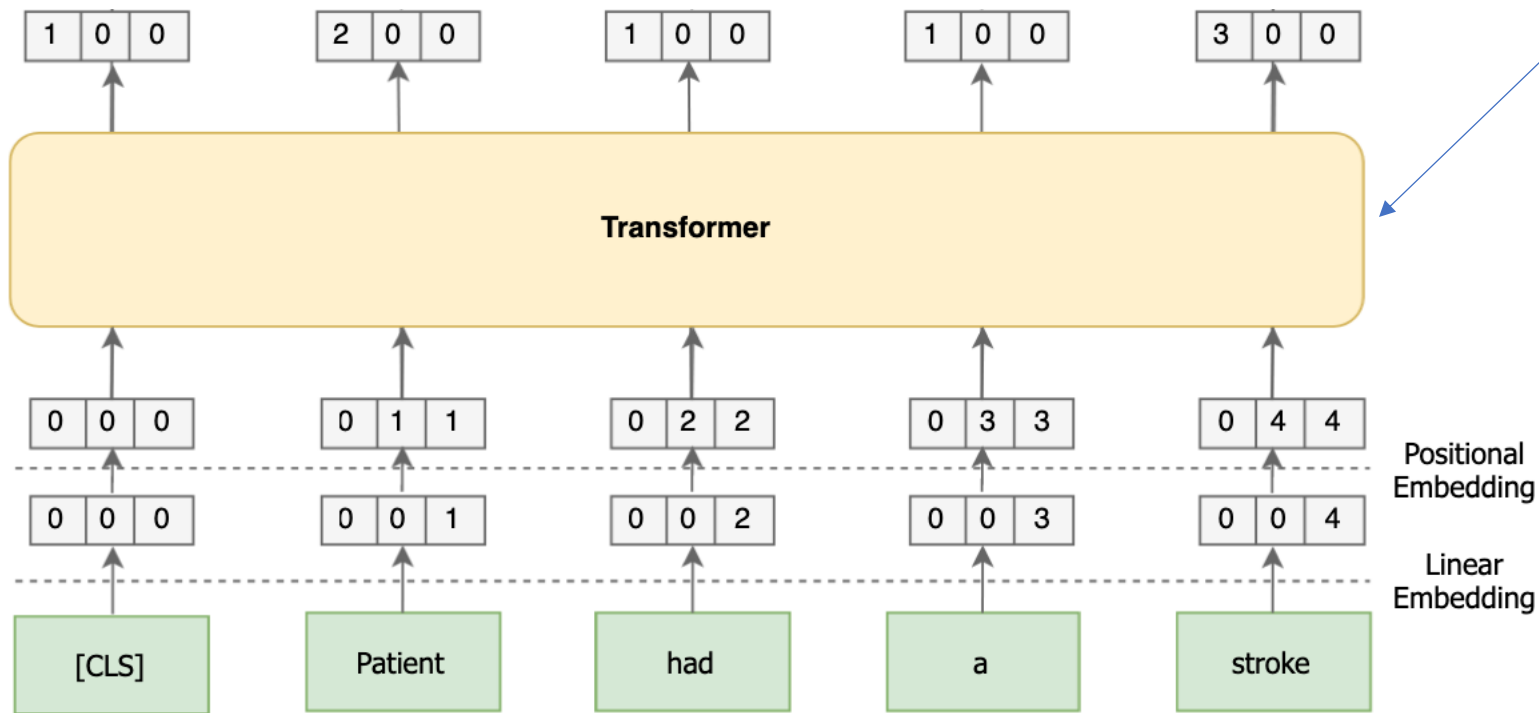


Transformer

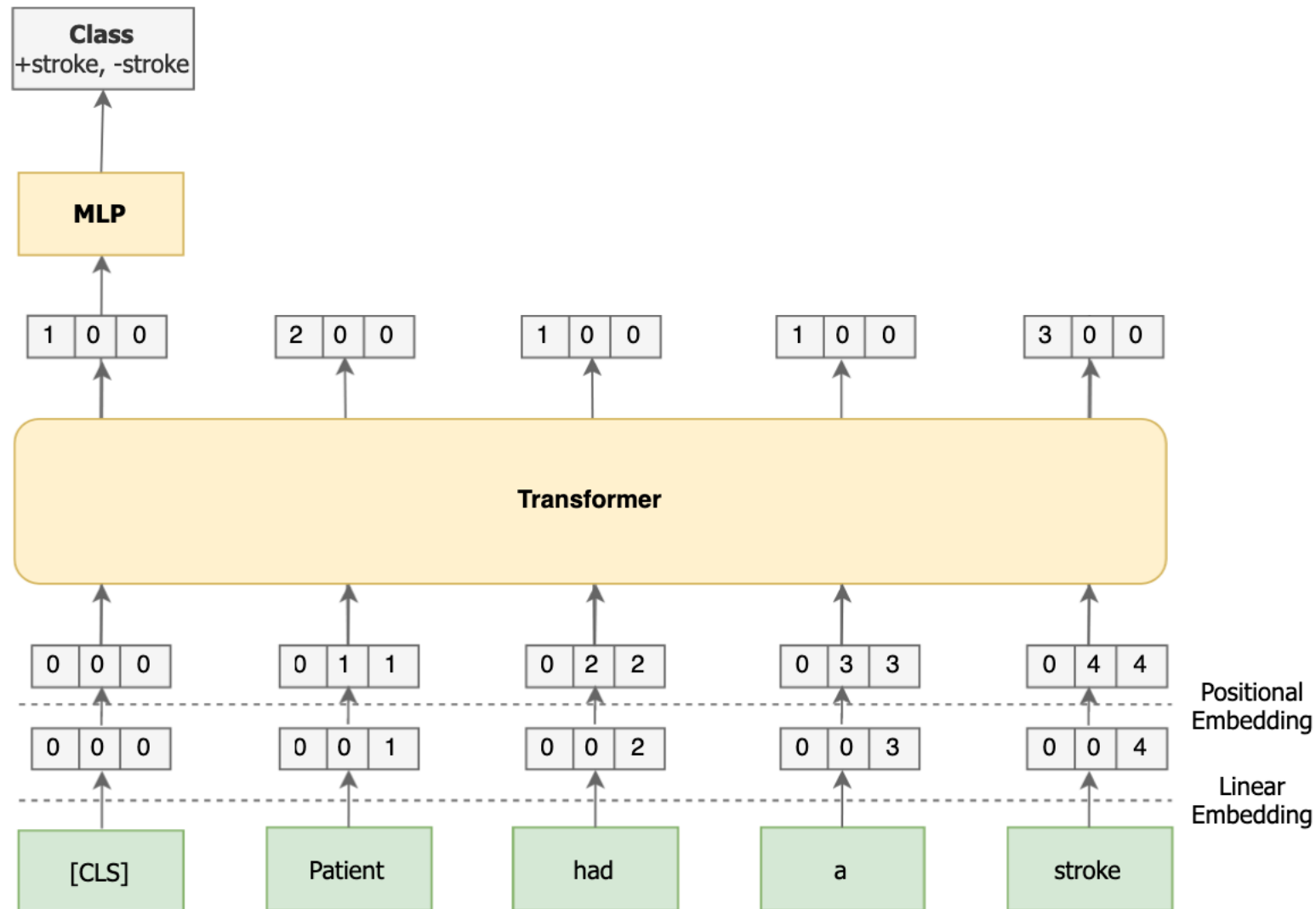


Transformer

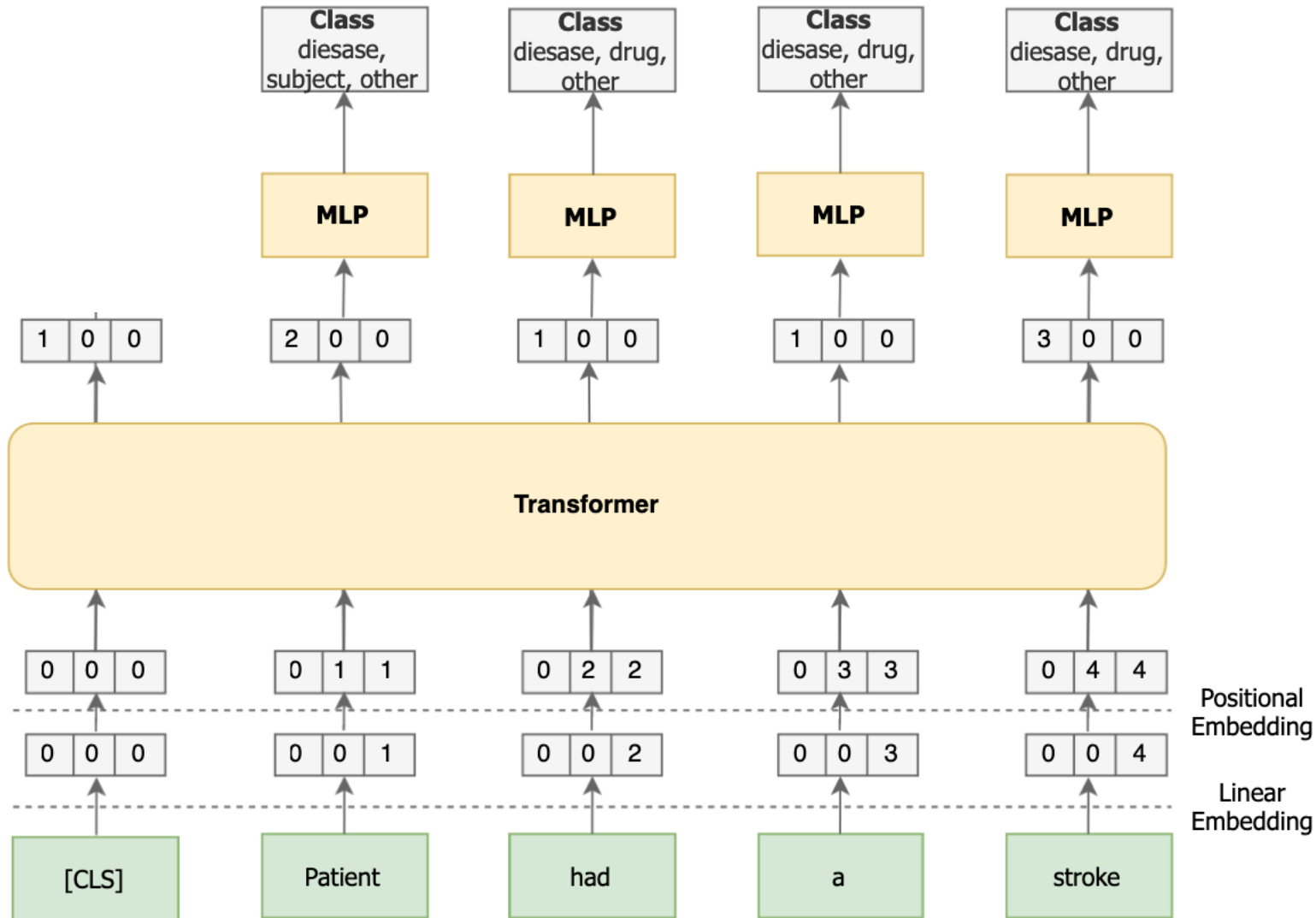
If you want to peek inside
the yellow box this is a
[great blog post](#)



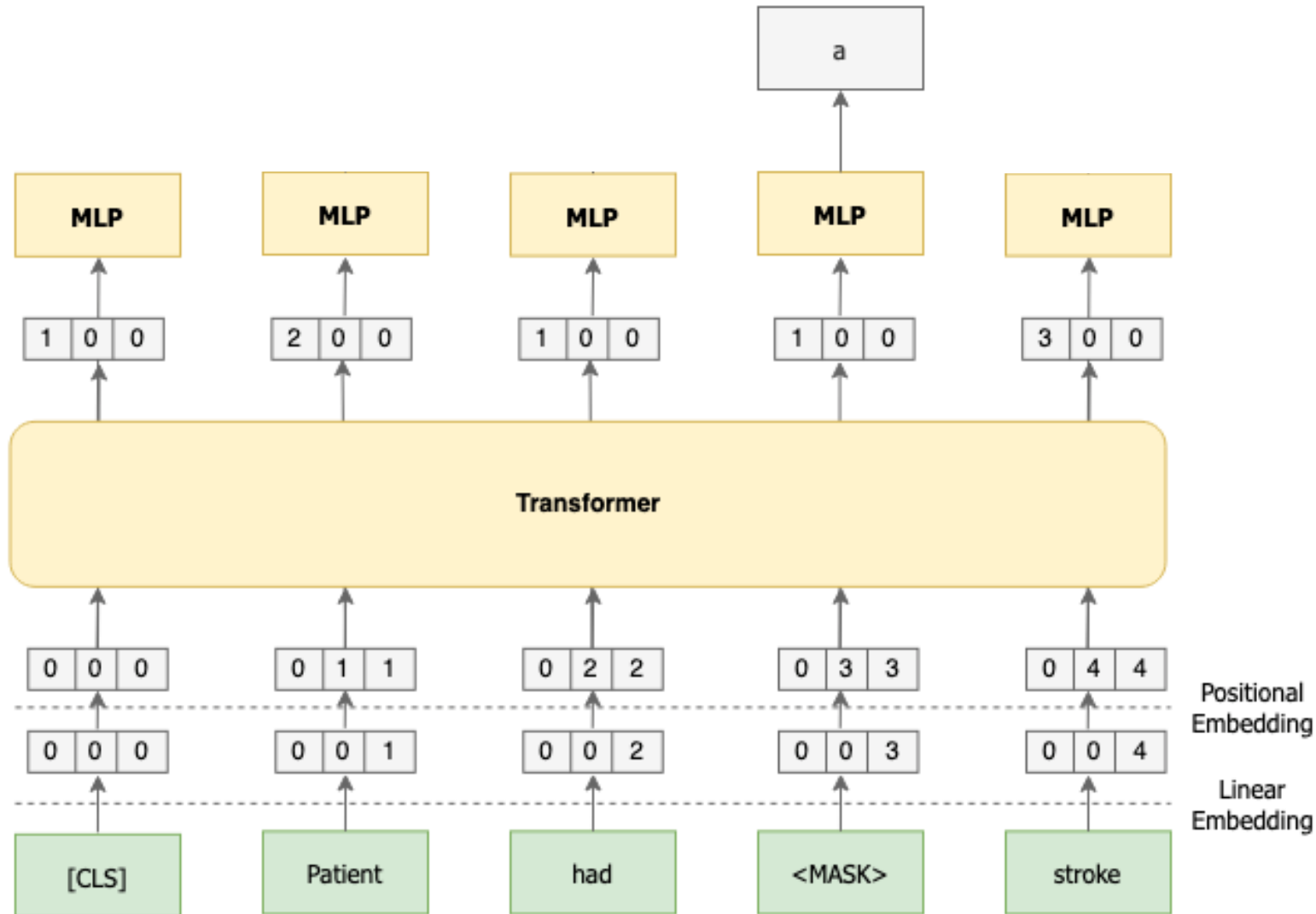
Supervised Learning: Classification



Supervised Learning : Named Entity Recognition



Pretraining: Masked Language Modelling



Pretraining: BERT SOTA ~202

bert-base-uncased like 1.35k

Fill-Mask Transformers PyTorch TensorFlow JAX Rust Core ML ONNX Safetensors bookcorpus wikipedia English bert exbert Inference Endpoints

arxiv:1810.04805 License: apache-2.0

Model card Files and versions Community 55

BERT base model (uncased)

Pretrained model on English language using a masked language modeling (MLM) objective. It was introduced in [this paper](#) and first released in [this repository](#). This model is uncased: it does not make a difference between english and English.

Downloads last month **39,669,693**

Safetensors Model size 110M params Tensor type F32

“For the pre-training corpus we use the BooksCorpus (800M words) (Zhu et al.,2015) and English Wikipedia (2,500M words).”

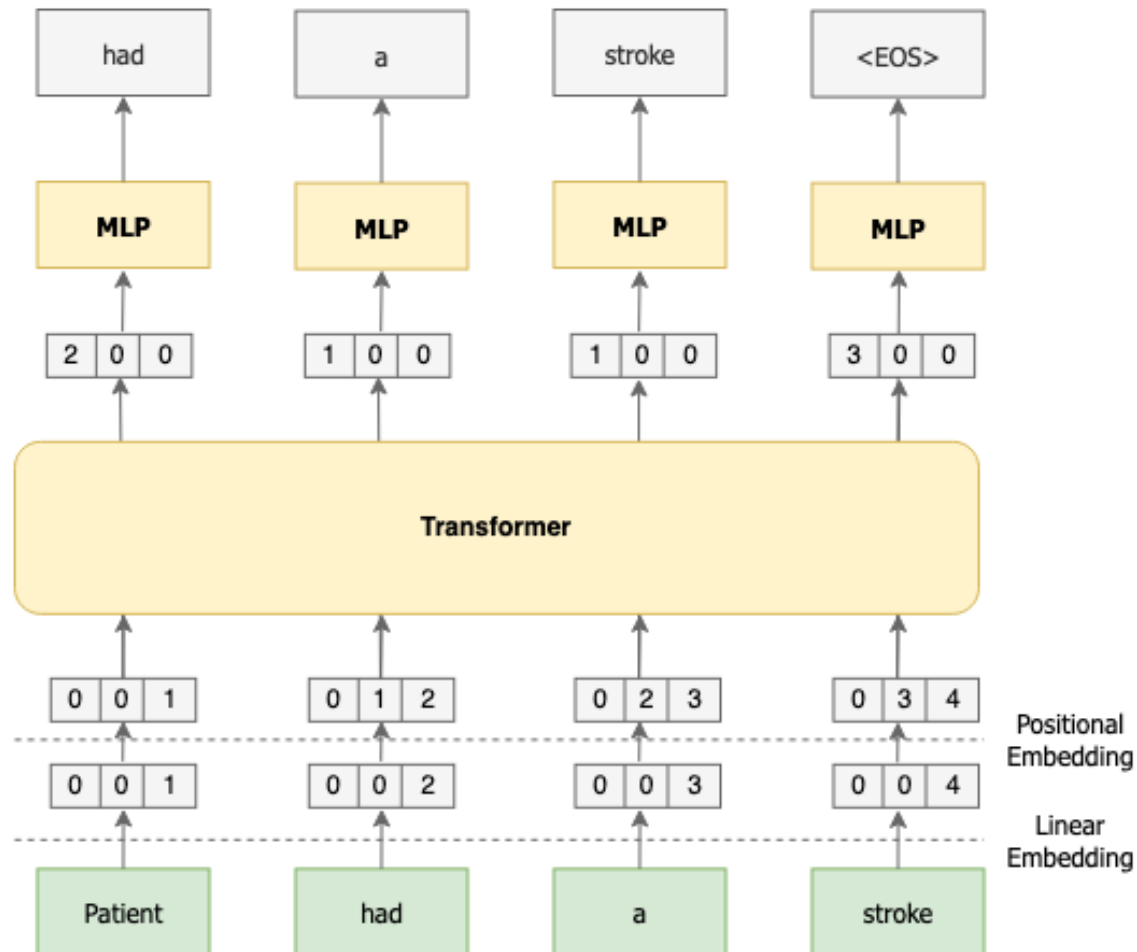
Supervised training for

- ICD 10 Coding
- Snomed concept extraction
- Adverse Drug Events

<https://huggingface.co/bert-base-uncased>

[BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding](#)

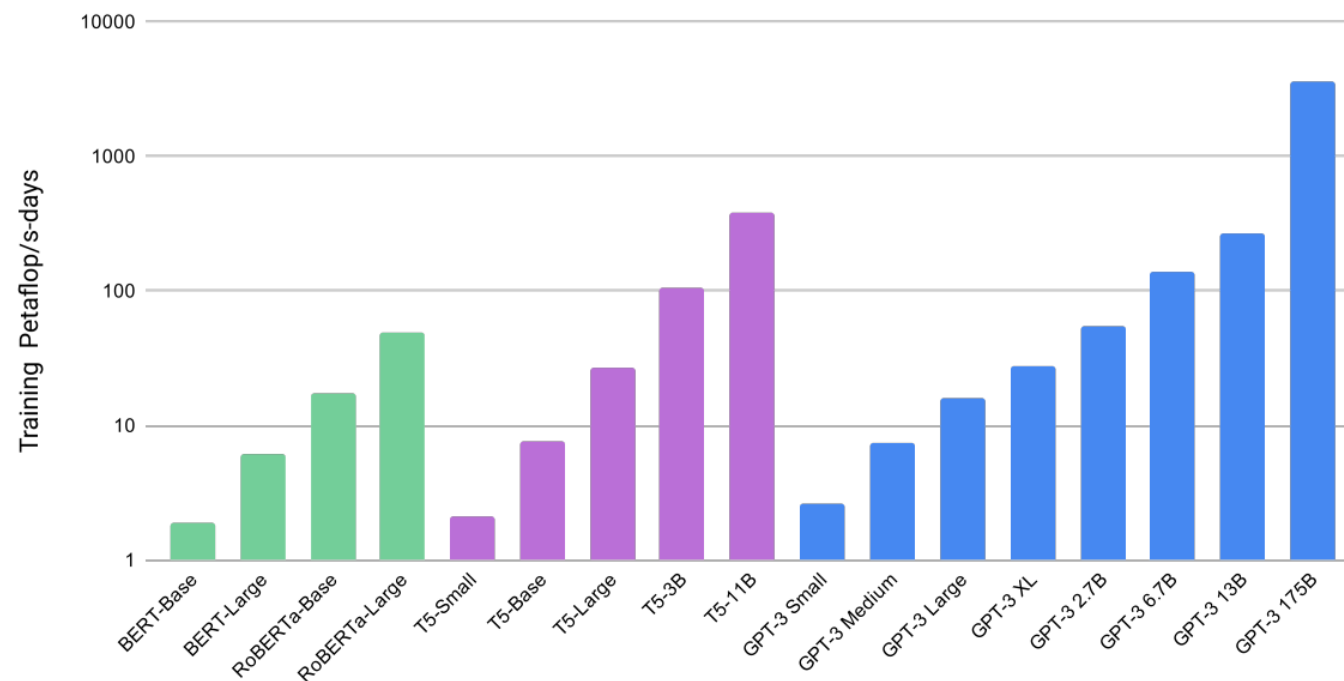
Next Token Prediction: BERT->GPT



*different attention mechanism as well

BERT vs GPT-3

	BERT	LLM
Pre-training objective	Masked Language Modelling	Next Token Prediction
Number Training Tokens	3200 Million	300 Billion
Number of Parameter	340 Million	175 Billion



[Brown, Tom, et al. "Language models are few-shot learners." *Advances in neural information processing systems* 33 \(2020\): 1877-1901.](#)

[Devlin, Jacob, et al. "Bert: Pre-training of deep bidirectional transformers for language understanding." *arXiv preprint arXiv:1810.04805* \(2018\).](#)

BERT vs GPT-3

Language Models are Few-Shot Learners

Tom B. Brown*	Benjamin Mann*	Nick Ryder*	Melanie Subbiah*	
Jared Kaplan [†]	Prafulla Dhariwal	Arvind Neelakantan	Pranav Shyam	Girish Sastry
Amanda Askell	Sandhini Agarwal	Ariel Herbert-Voss	Gretchen Krueger	Tom Henighan
Rewon Child	Aditya Ramesh	Daniel M. Ziegler	Jeffrey Wu	Clemens Winter
Christopher Hesse	Mark Chen	Eric Sigler	Mateusz Litwin	Scott Gray
Benjamin Chess	Jack Clark	Christopher Berner		
Sam McCandlish	Alec Radford	Ilya Sutskever	Dario Amodei	

OpenAI

Abstract

3 Results

- 3.1 Language Modeling, Cloze, and Completion Tasks
- 3.2 Closed Book Question Answering
- 3.3 Translation
- 3.4 Winograd-Style Tasks
- 3.5 Common Sense Reasoning
- 3.6 Reading Comprehension
- 3.7 SuperGLUE
- 3.8 NLI
- 3.9 Synthetic and Qualitative Tasks

Scaling Laws

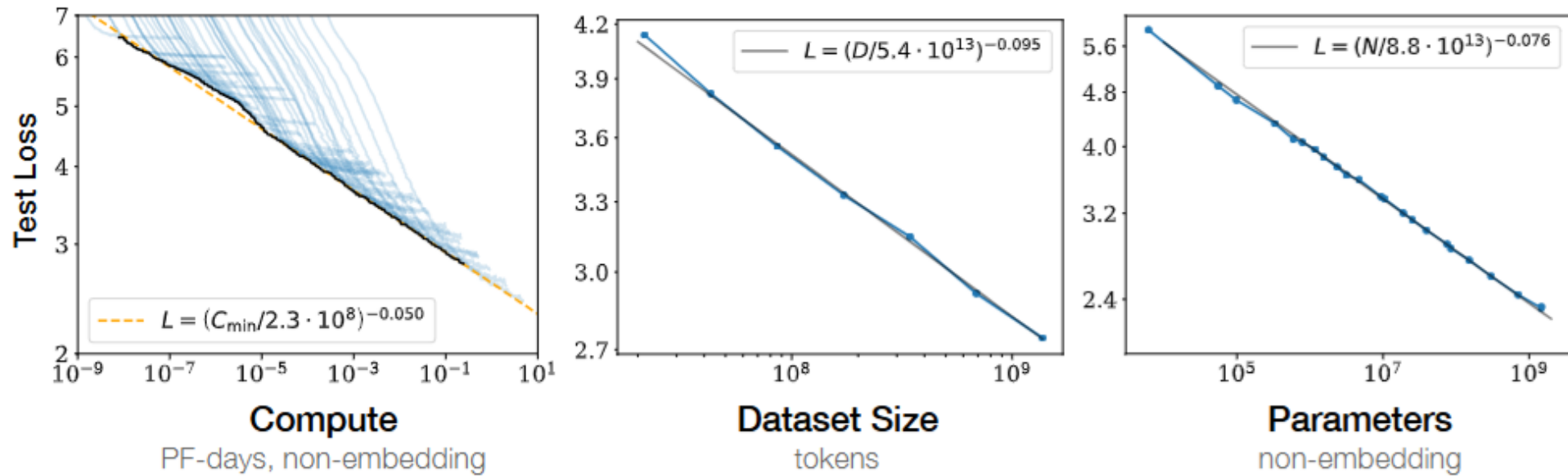
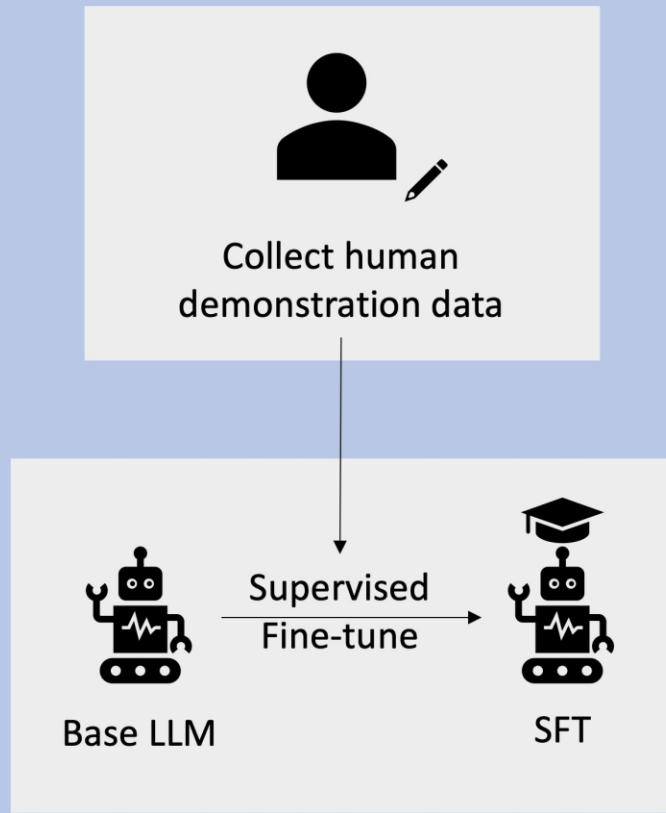


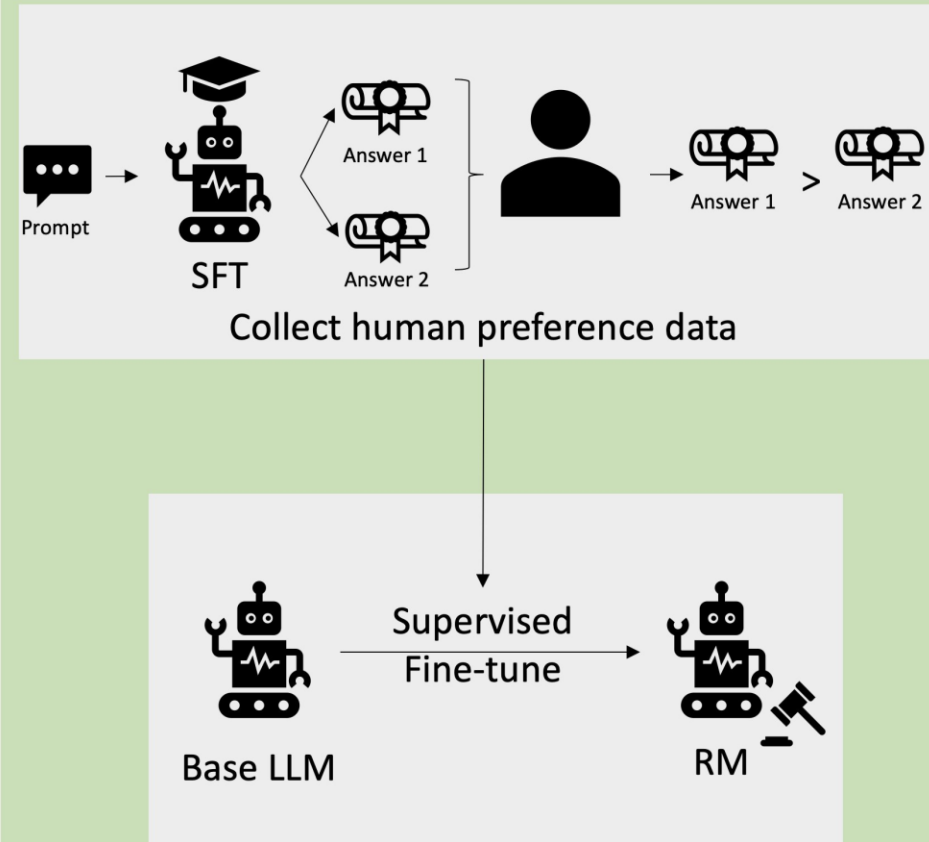
Figure 1 Language modeling performance improves smoothly as we increase the model size, dataset size, and amount of compute² used for training. For optimal performance all three factors must be scaled up in tandem. Empirical performance has a power-law relationship with each individual factor when not bottlenecked by the other two.

RLHF: GPT-> ChatGPT

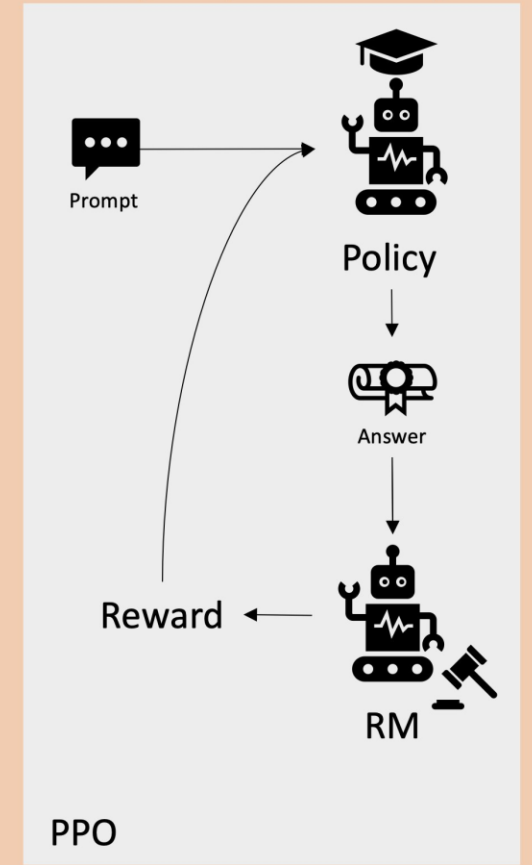
Step 1 Supervised Fine-Tuning



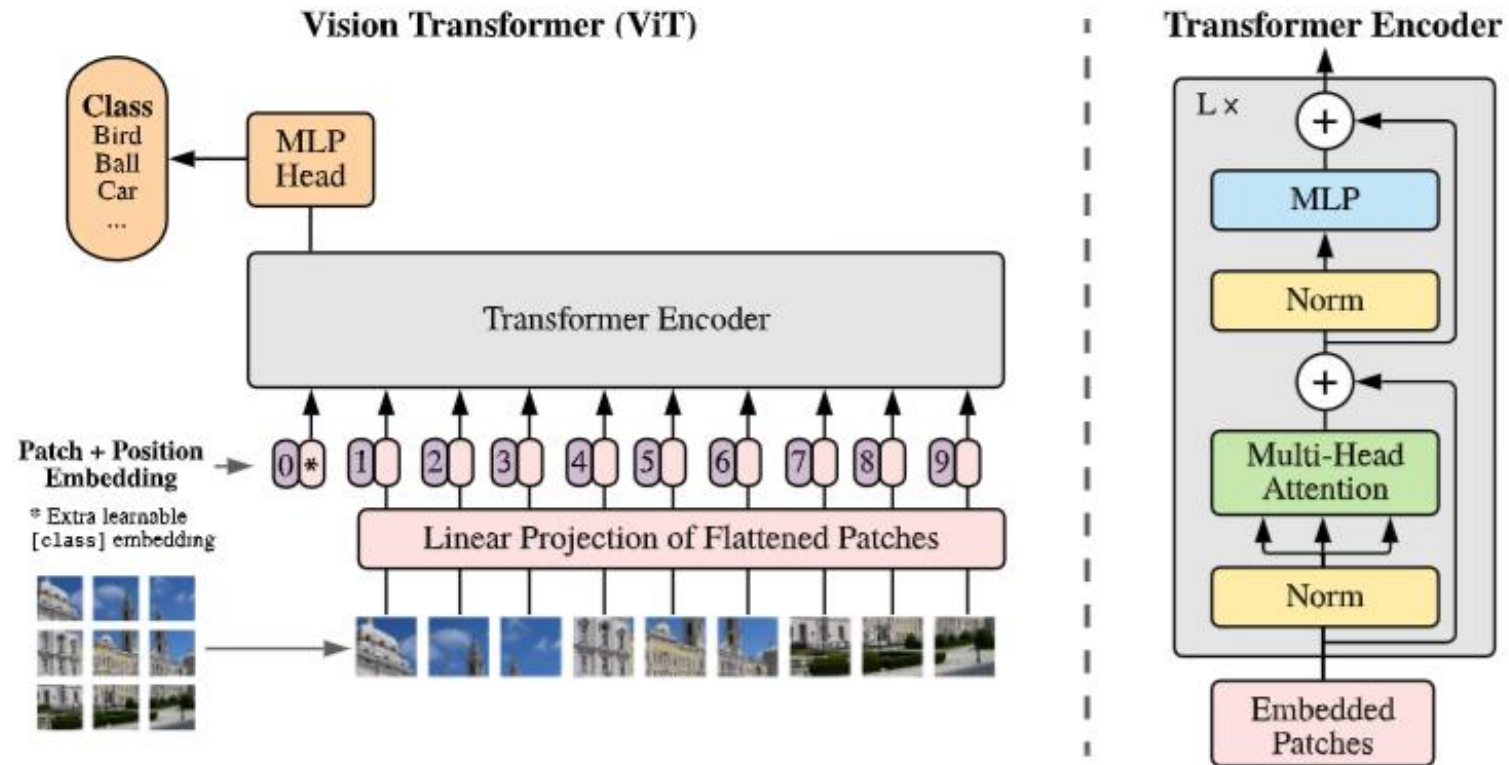
Step 2 Training a Reward Model



Step 3 Optimize Policy



Multi-Modal: Will it tokenize?



What can I help with?

Write the answer to Prof Taylor's latest assignment



Search



Reason



Create image



Code



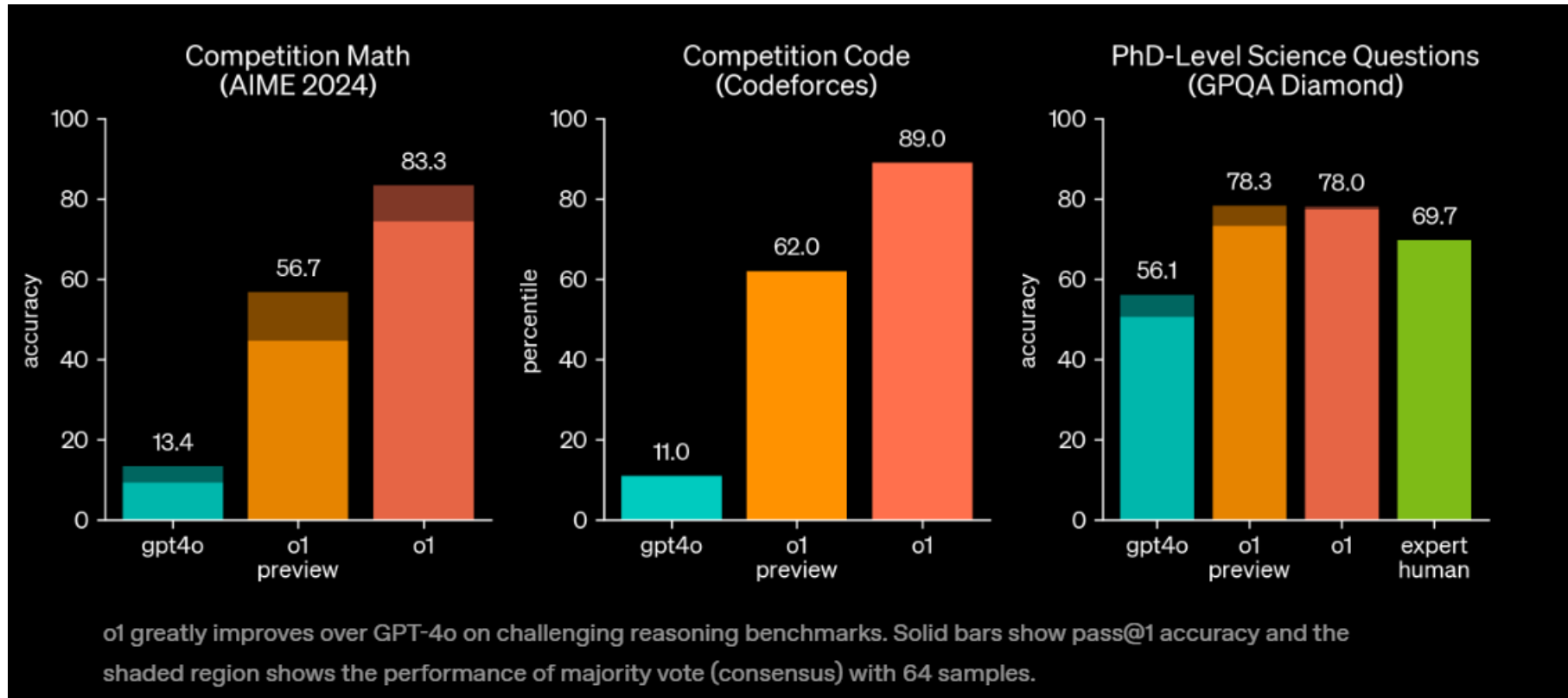
Summarize text



Help me write

More

Reasoning Models



Reasoning Models: DeepSeek



Chain of Thought

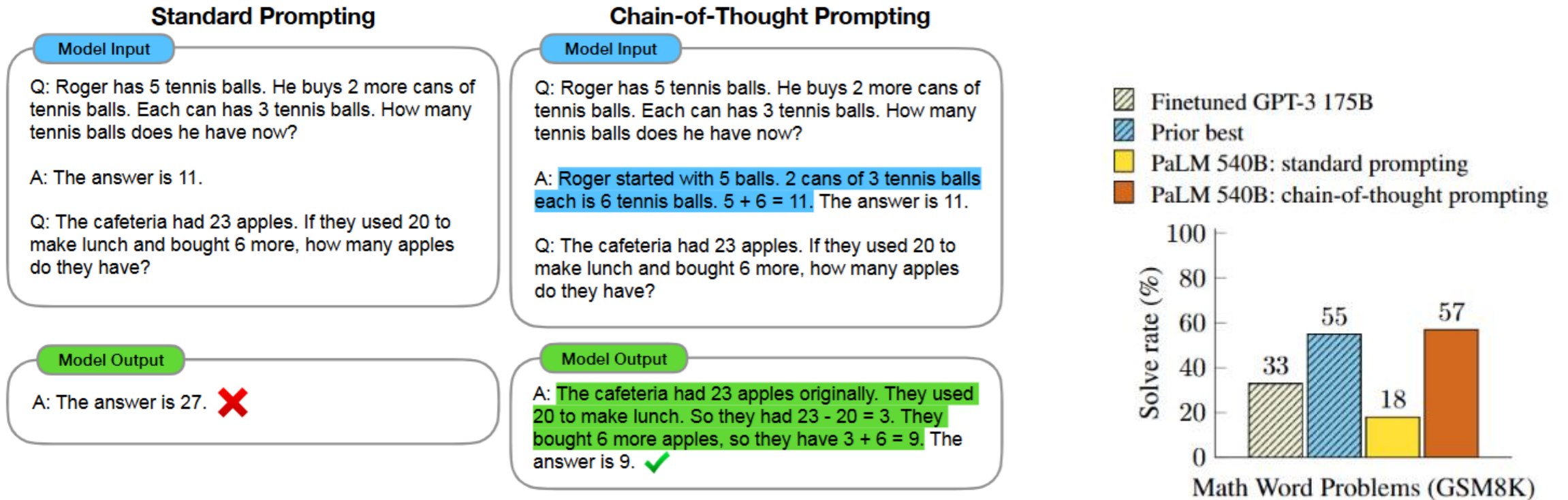
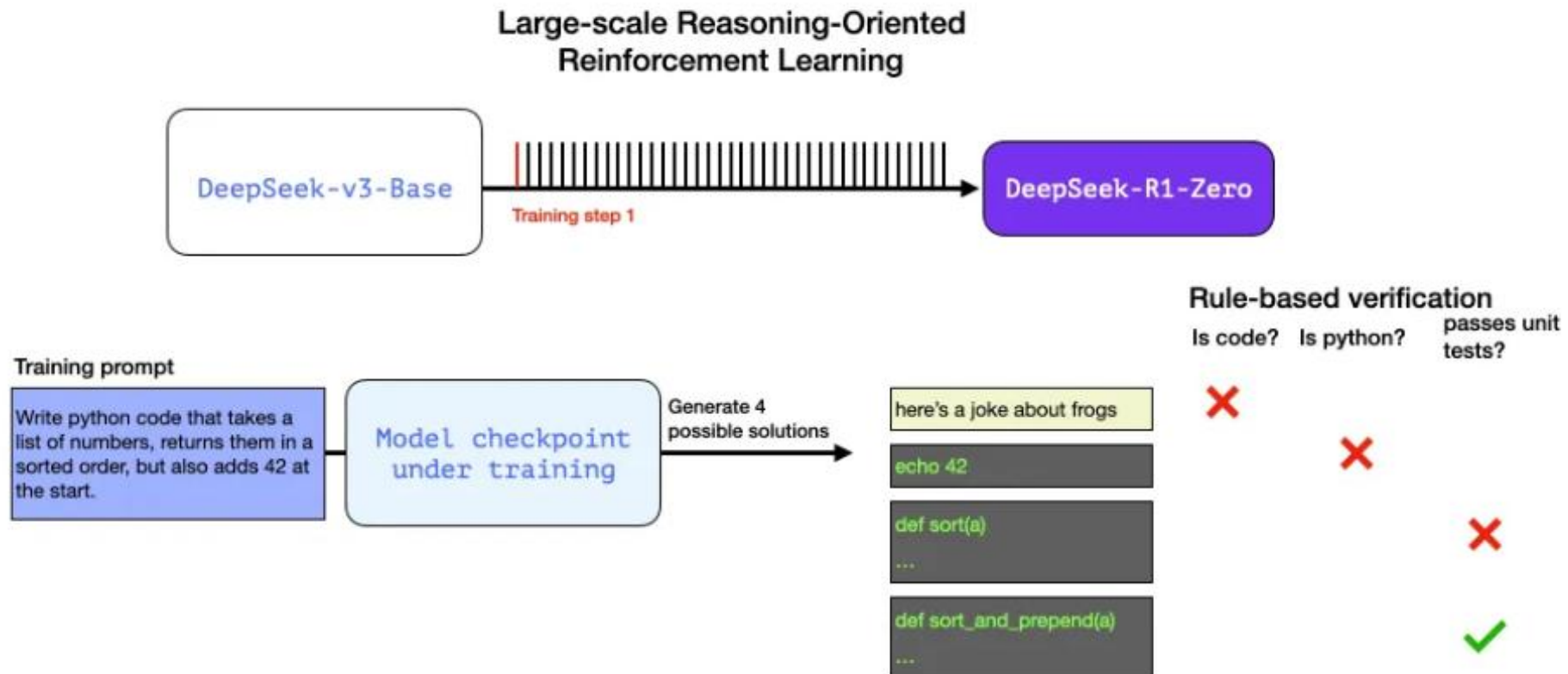


Figure 1: Chain-of-thought prompting enables large language models to tackle complex arithmetic, commonsense, and symbolic reasoning tasks. Chain-of-thought reasoning processes are highlighted.

Reasoning Models: DeepSeek



Reasoning Models: DeepSeek

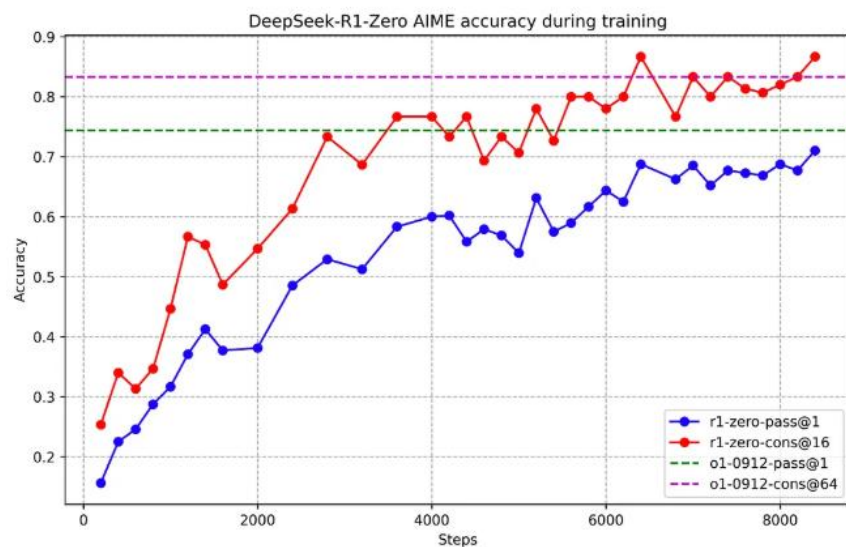


Figure 2 | AIME accuracy of DeepSeek-R1-Zero during training. For each question, we sample 16 responses and calculate the overall average accuracy to ensure a stable evaluation.

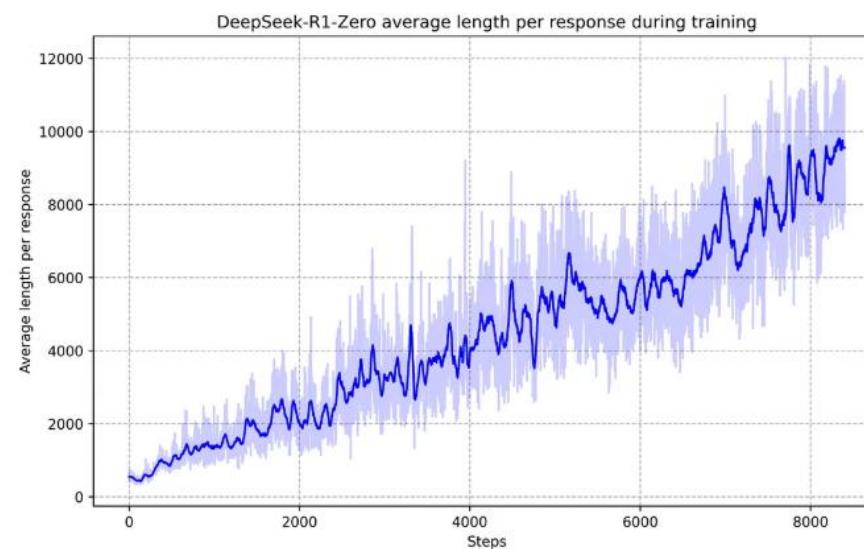
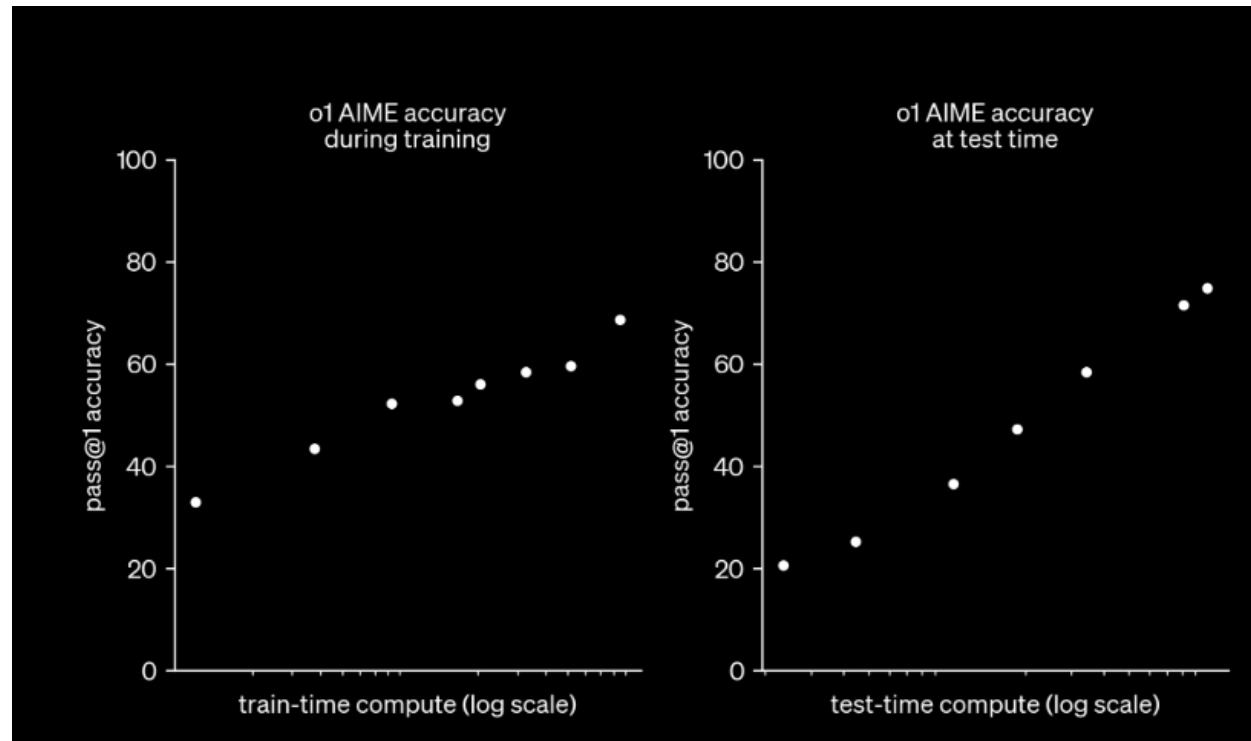
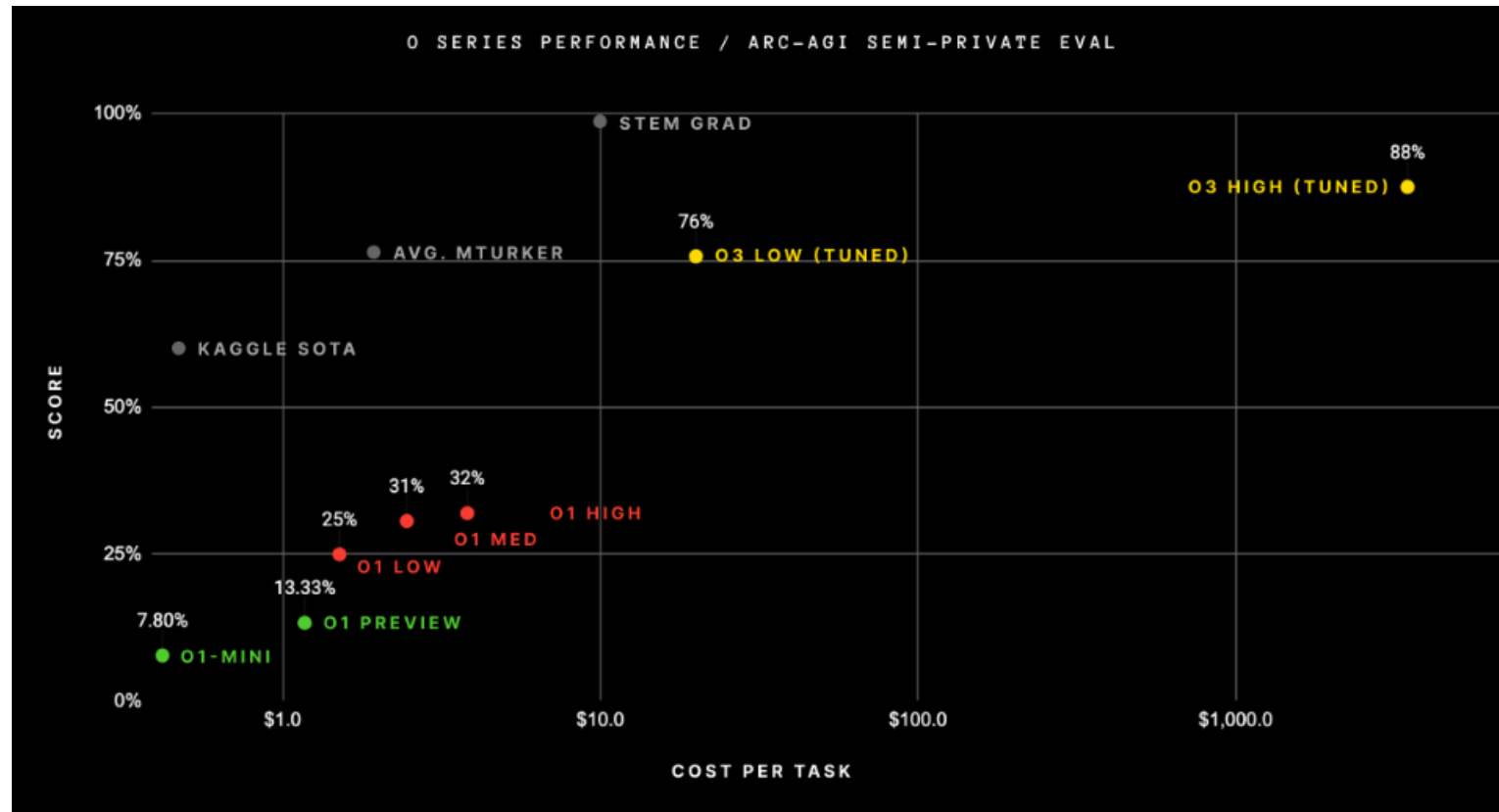


Figure 3 | The average response length of DeepSeek-R1-Zero on the training set during the RL process. DeepSeek-R1-Zero naturally learns to solve reasoning tasks with more thinking time.

Reasoning Models: Test time compute



Reasoning Models: o3?



Take away points

- How to make ChatGPT. Scale all of:
 1. Transformer model
 2. Large-scale unsupervised pre-training
 1. Multi-modal: Will it tokenize?
 3. RLHF to make chatty

Gives a general few/zero-shot model
- Reasoning models
 - Seems to be even more RL(HF)
 - Watch this space...

I just need
the main ideas



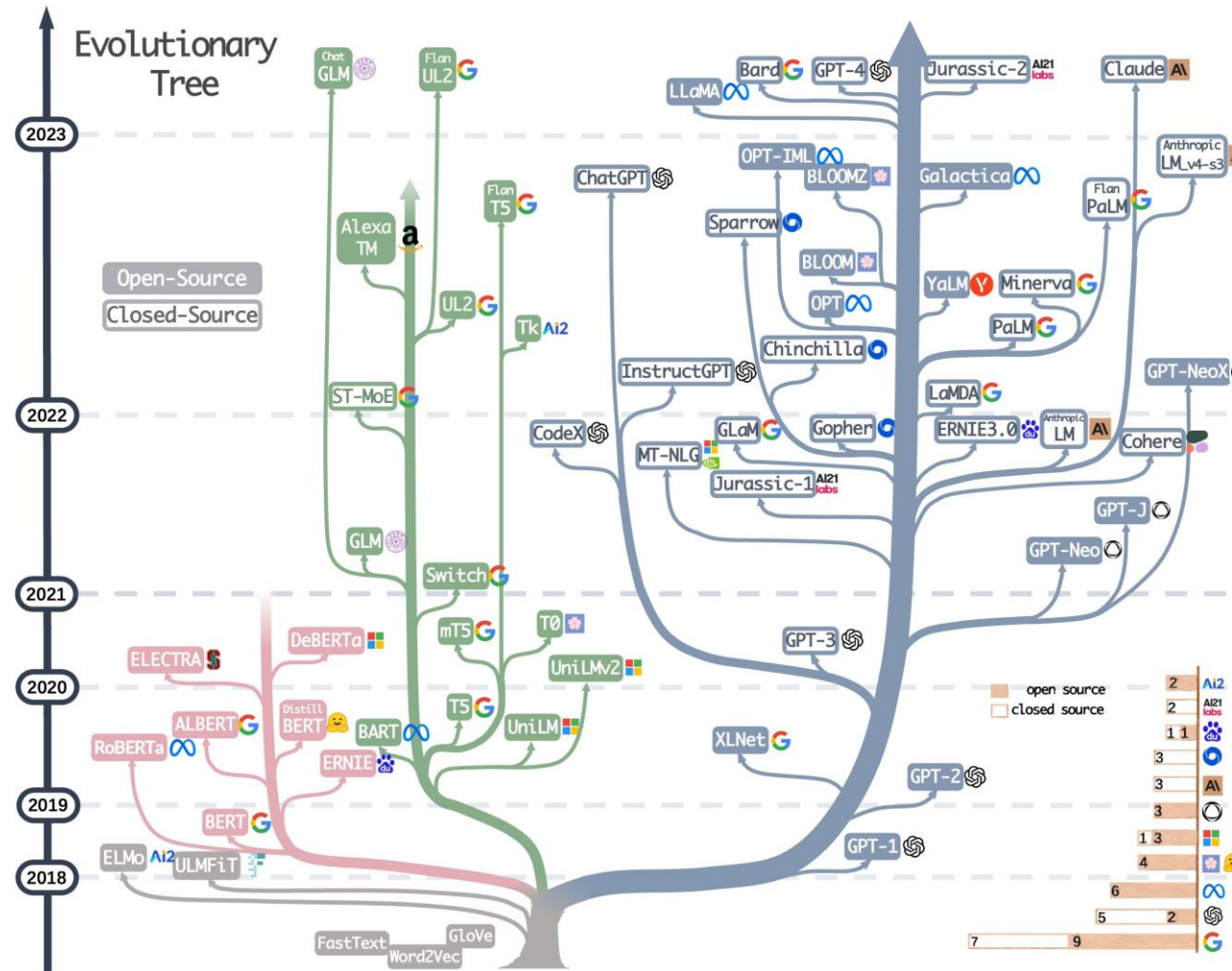




Have a break, have a Kit Kat.®

Part 2- How to use an LLM

What model?



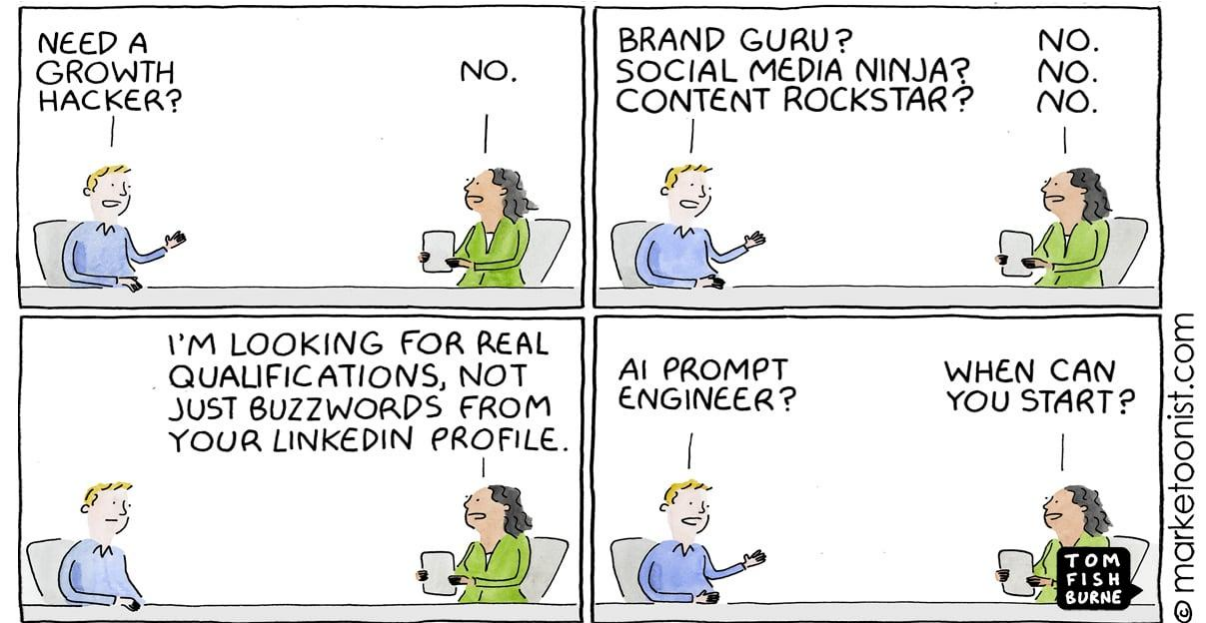
Open Source/Open Weight

- Open Weights = Model weights are freely accessible on the internet
- Open Source = Training code, data and model are freely accessible on the internet
- Current ~6-12 month lag with closed source
- Can be finetuned using supervised learning by approximation (e.g. LORA)
- 3rd Party vs Local Hosting \$\$\$
 - “Small” large language models



Prompt Engineering

- Be explicit in input/output
- Few shot prompting
- Chain of thought
- Self-consistency
- Structured Output



Prompt Engineering

System:

""You are a consultant doctor tasked with writing a patients discharge summary.

A user will provide you with a list of clinical notes from a hospital stay from which you will write a discharge summary.

Each clinical note has a title of the format [Title]: [timestamp year-month-day hour:min].

Clinical notes are ordered by ascending timestamp.

Only the information in the clinical notes provided by the most recent user message can be used for this task.

The discharge summary must be written in accordance with the following json schema.

{json_schema}

All fields are required.

If the relevant information is not present in the clinical notes, fields can be filled with an empty string or list.

Expand all acronyms to their full terms.""

User:

""Clinical Notes

{rcp_exemplar_clinican_notes}

Please write a discharge summary only using the information in this message's clinical notes.

The discharge summary must be written in accordance with the json schema given in the system message.""

Assistant:

""{rcp_exemplar_discharge_summary_json}""

One-shot
example

User:

""Clinical Notes

{input_clinician_notes}

Please write a discharge summary only using the information in this message's clinical notes.

The discharge summary must be written in accordance with the json schema given in the system message.""

Assistant (excerpt):

""...

admission_details": {

"reason_for_admission": "Chest tightness pain, breathlessness, nausea and dizziness started at 6 am.",

"admission_method": "Emergency admission via London Ambulance Service",

"relevant_past_medical_and_mental_health_history": [

"Type 2 Diabetes medication (tablets)",

"Hypertension",

"Chronic Obstructive Pulmonary Disease"

]

},...""

Question Answering on MedQA

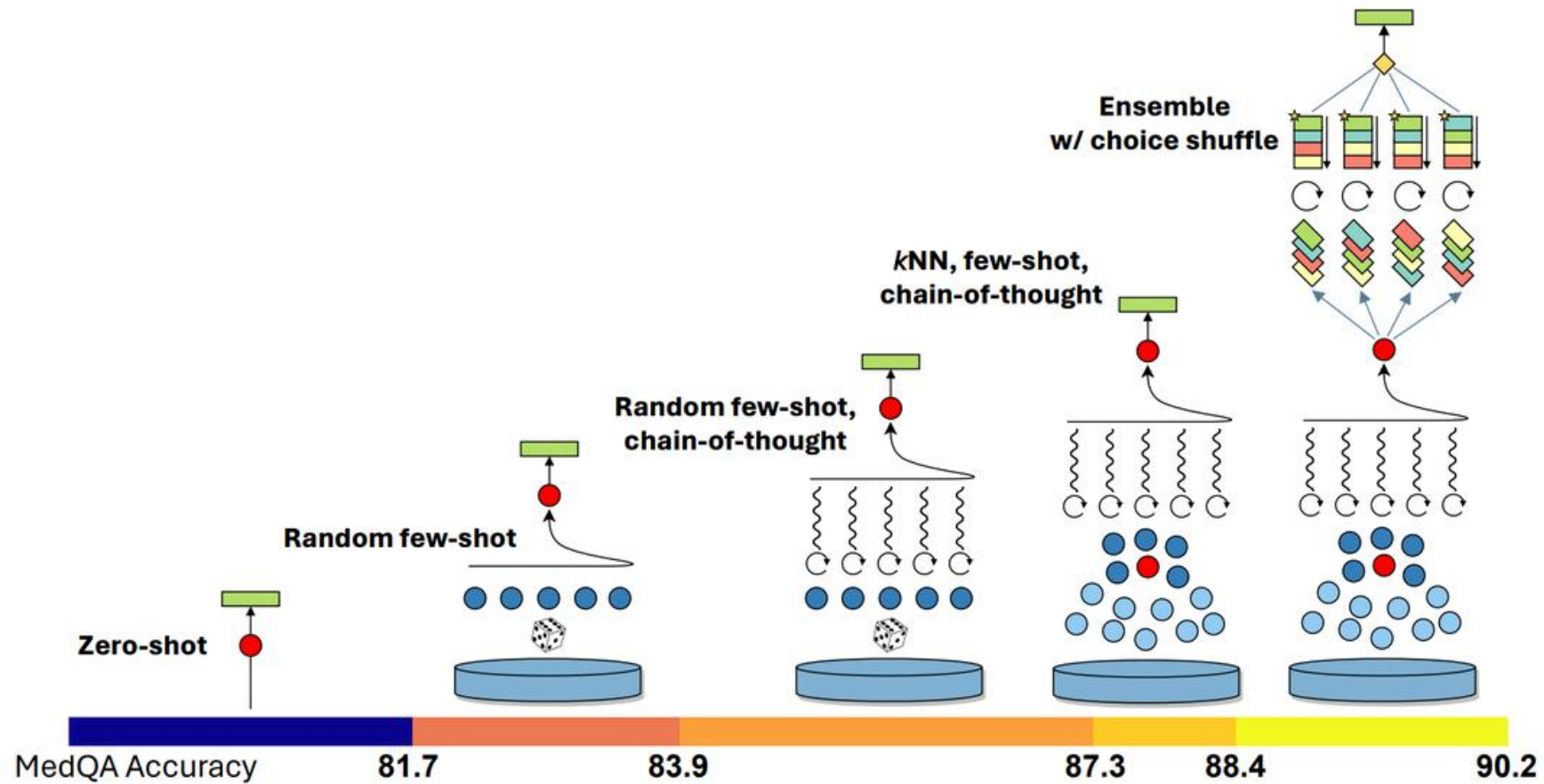
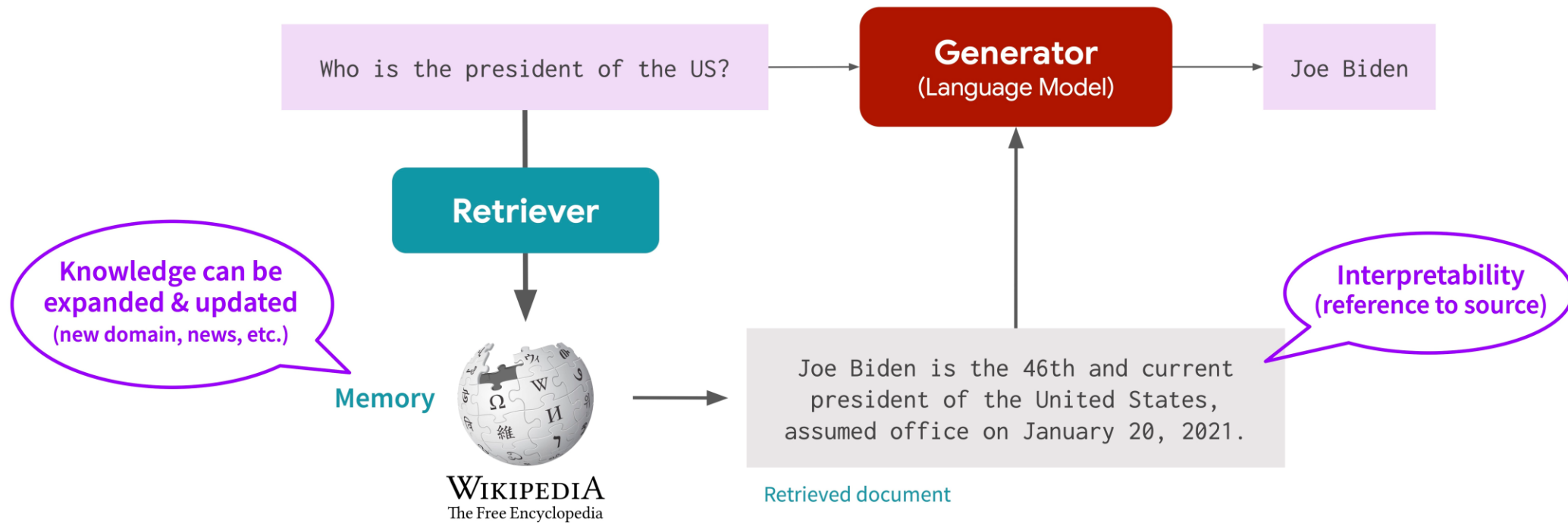


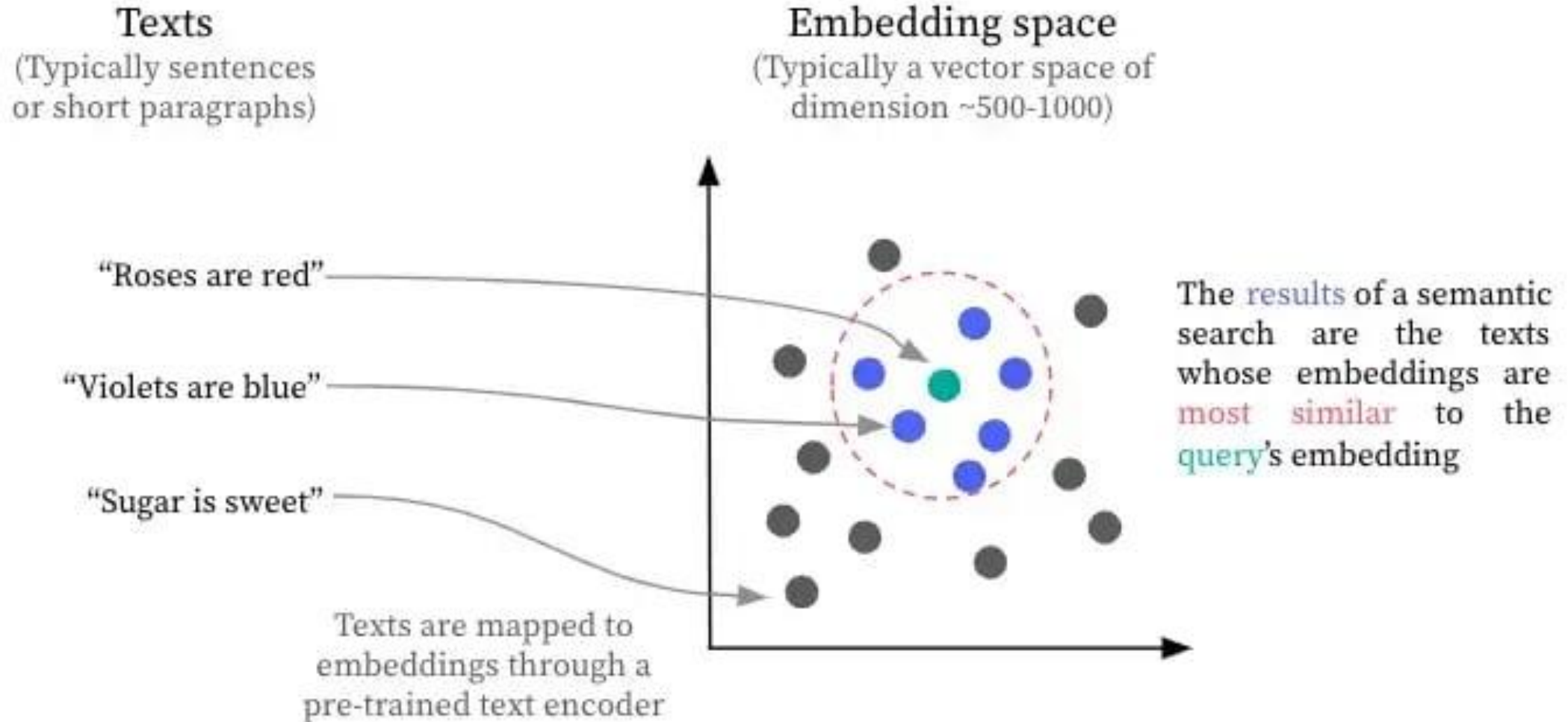
Figure 1: Visual illustration of Medprompt components and additive contributions to performance on the MedQA benchmark. Prompting strategy combines kNN-based few-shot example selection, GPT-4-generated chain-of-thought prompting, and answer-choice shuffled ensembling.

Retrieval Augmented Generation

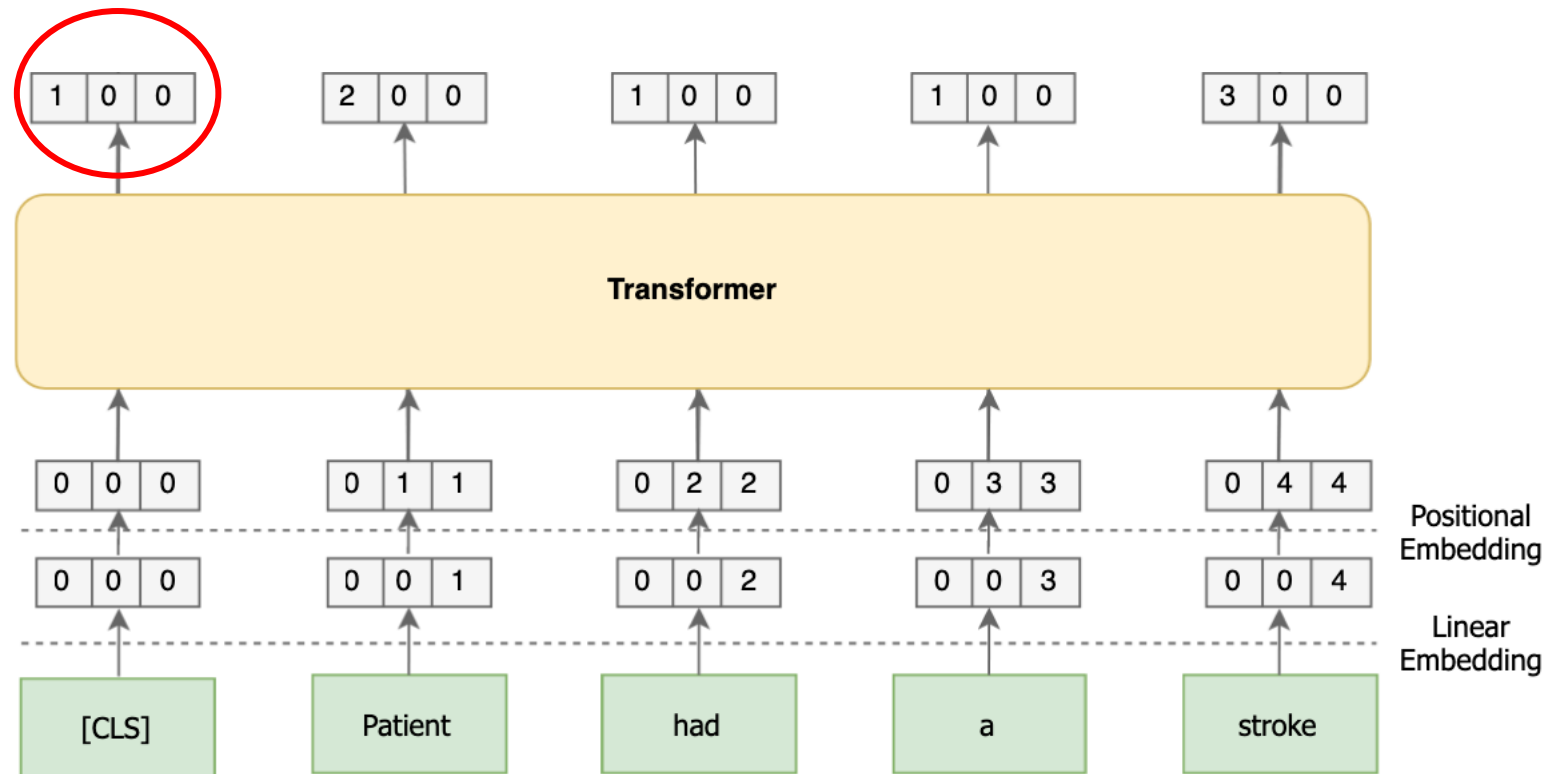
Retrieval augmentation



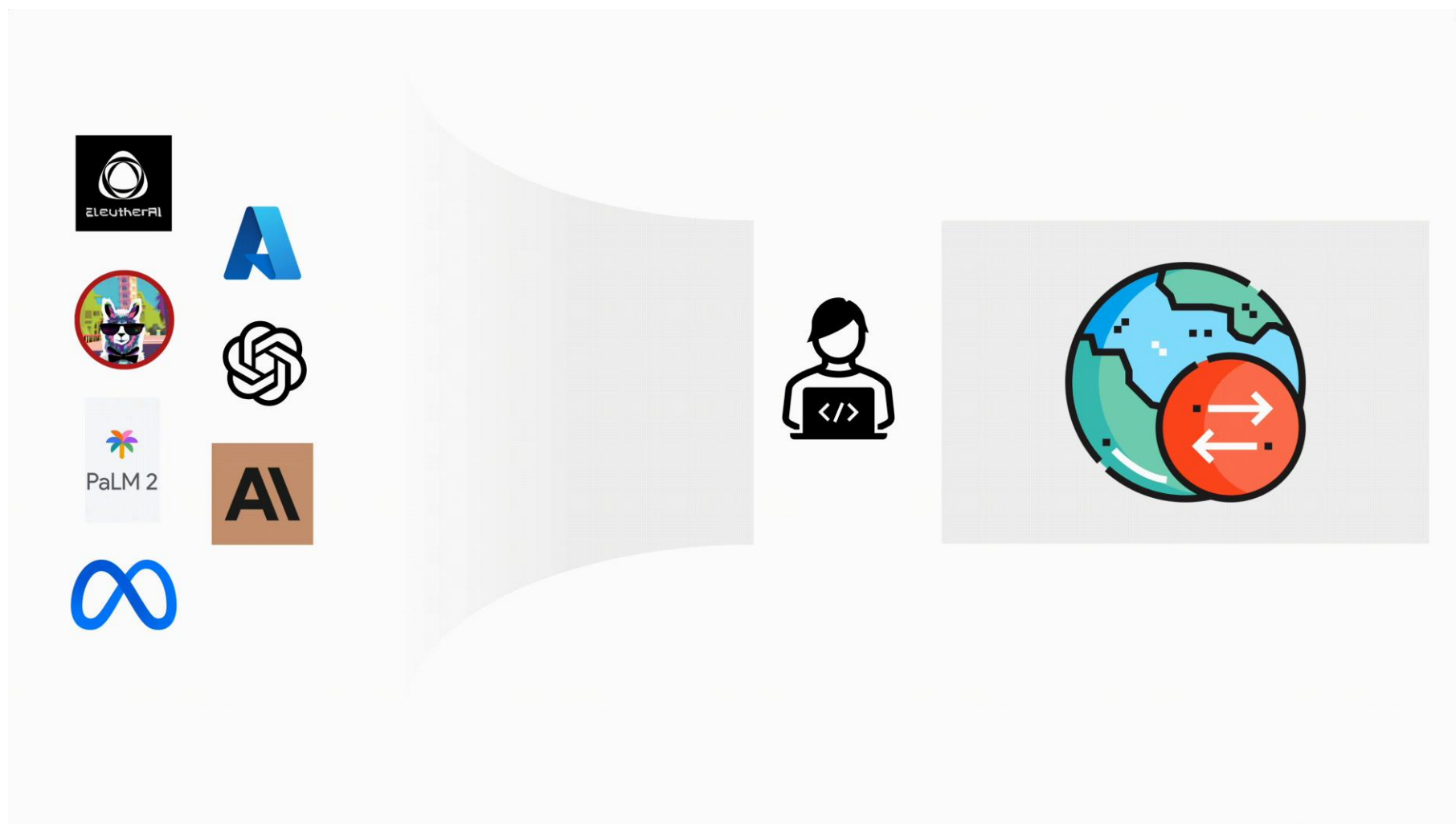
Semantic Search



Text -> Vector Embeddings



Agents



Recap

- What Model?
 - Closed vs open source “frontier models”
 - 3rd party vs local hosting
- Prompt engineering
 - An ‘art not a science’
 - Can make an annoyingly big difference...
- Add own data using RAG

I just need
the main ideas





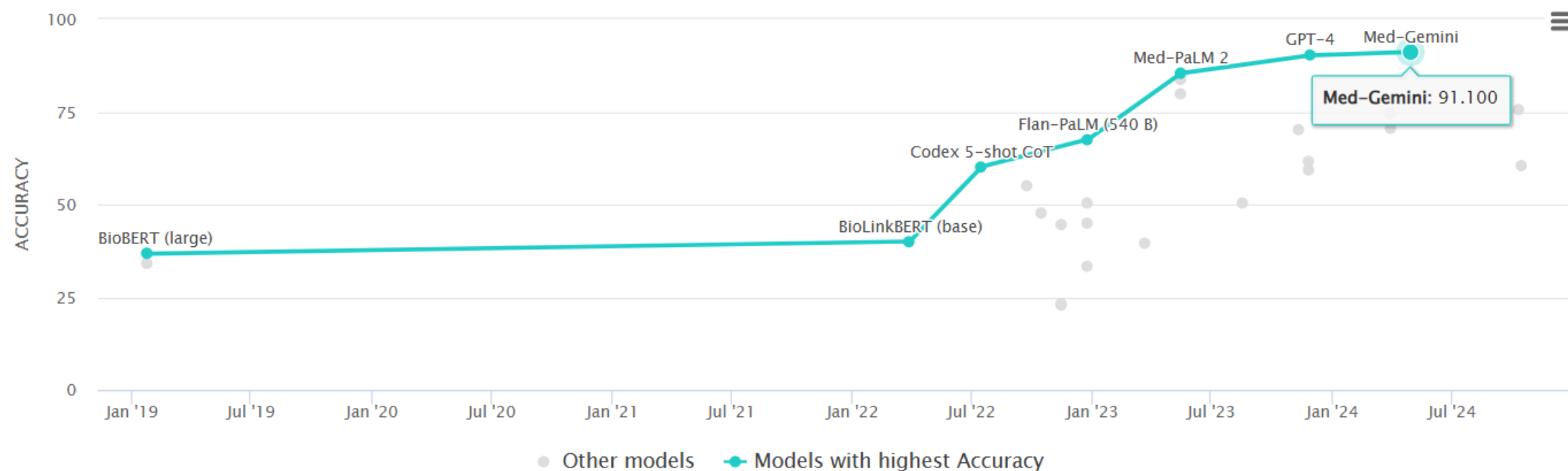
What does this mean for
medicine?

Question Answering on MedQA

Leaderboard

Dataset

View Accuracy by Date for All models



Automated Diagnoses

Read the case below and answer the question provided after the case.

Format your response in markdown syntax to create paragraphs and bullet points. Use '

' to start a new paragraph. Each paragraph should be 100 words or less. Use bullet points to list multiple options. Use '
*' to start a new bullet point. Emphasize important phrases like headlines. Use '**' right before and right after a phrase to emphasize it. There must be NO space in between '**' and the phrase you try to emphasize.

Case: [Case Text]

Question (suggested initial question is "What are the top 10 most likely diagnoses and why (be precise?)"): [Question]

Answer:

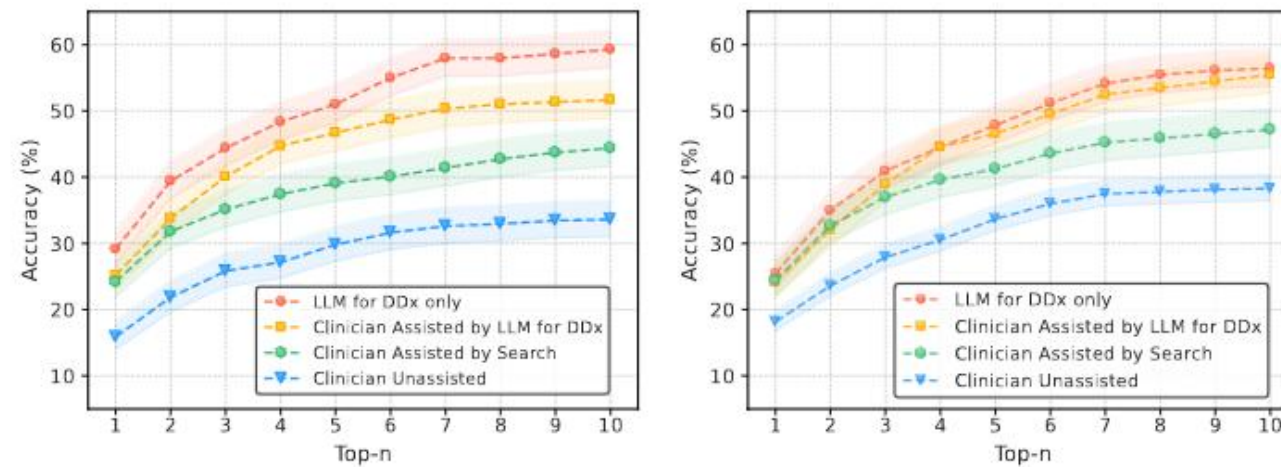


Figure 5 | Top-n Accuracy. (left) The percentage of DDx lists with the final diagnosis through human evaluation. (right) The percentage of DDx lists with the final diagnosis through automated evaluation.

Predictive Modelling

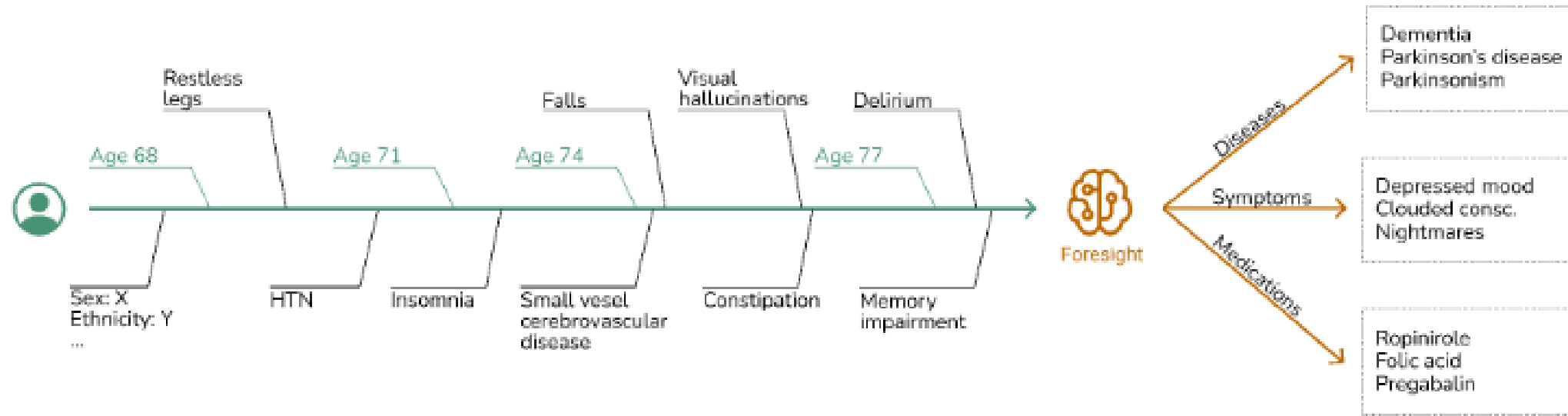


Figure 2. The left portion of the timeline represents the existing/historical data for a patient and the right portion are forecasts from Foresight for different biomedical concept types.

Automating Bureaucracy

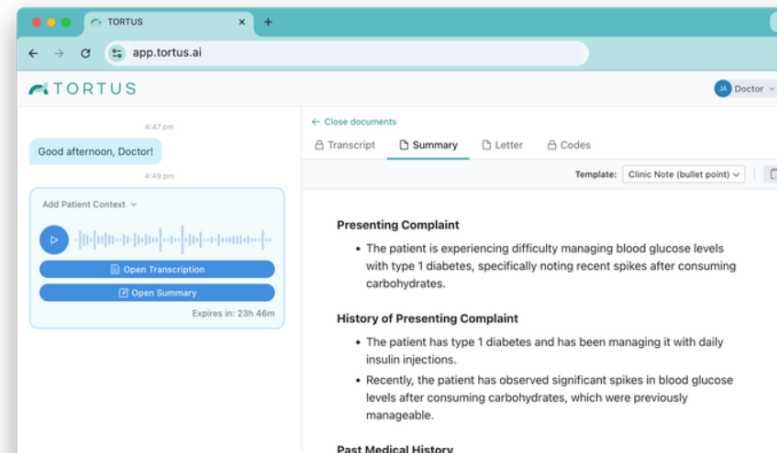
GOSH pilots AI tool to give clinicians more quality-time with patients

11 Nov 2024, 7 a.m.

A hands-free approach to drafting clinic notes

Typically, during outpatient consultations with patients, clinicians will dictate or manually type notes into the computer and compose letters during and after the consultation. The TORTUS assistant automates this process, **using ambient voice technology with generative AI to listen to the consultation and draft a clinic note and letter**, which are then edited and authorised by the clinician before being uploaded to the electronic health record system and sent to patients and their families.

The technology means clinicians can take a more 'hands-free' approach to drafting notes and follow-up letters as TORTUS does it for them, leaving them able to spend more time on direct patient care during outpatient appointments. In early testing in simulated clinics, all clinicians agreed that the AI helped them give their full attention to their patients when using the tool, without decreasing the quality of the clinic note or letter.



Automating Bureaucracy



Software

About Us ▾

Newsroom

[← Back to Software](#)

Artificial Intelligence

AI is improving healthcare and has great potential to do much more. AI is used in many forms, including generative AI and predictive models, throughout Epic's software today with more than 100 significant development projects underway.

Here's a peek at what's available today, how health systems are using it, and what's in store.

We encourage Epic community members to reach out to their BFFs to learn more.



Clinician Efficiency

[details →](#)



Insights from Data

[details →](#)



Revenue Cycle

[details →](#)

Research Proof of Concept -> Bedside

- Data governance
- Compute requirements
- Regulation
- Robust evaluation
- How to show ROI -> health economics/patient outcomes

WE NEED YOU!



Recap

- Medical use cases
 - Benchmarks being saturated
 - Automating bureaucracy 'low hanging fruit'
 - Commercialised ASAP
- Barriers
 - Robust evaluation
 - Getting hospitals 'AI-ready'
 - Regulation?

I just need
the main ideas



Buzzwords

- LLM
- GPT
- Transformer
- Next token prediction
- Few shot prompting
- Chain of Thought
- Self-Consistency
- Finetuning
- Multi Modal
- Foundation Models
- RLHF
- Scaling Laws
- Open Source
- RAG
- Semantic Search
- Embeddings
- Agents
- Reasoning Model

I just want to know more about...

- How transformers work
 - [Slightly more technical LLM overview talk](#)
 - [Great blog post](#)
 - [Original paper](#)
 - [Code GPT from scratch](#)
- How to make an LLMs-based “thing”
 - Play around in OpenAI/Azure Playgrounds
 - [Popular coding framework](#)
 - [Popular RAG framework](#)
- LLMs and Medicine
 - Links to all papers at bottom of slides
 - [Recent narrative review](#)



