

LLM, AI Agent and Vibe Coding

Introduction and Recipe

Yue Ma, Postdoctoral Scholar at HDSI, UCSD
06/22/2026 @ 0vββ AI Summer School

**A special acknowledgement
in my PhD defense slides...**

Acknowledgement

Large Language Models



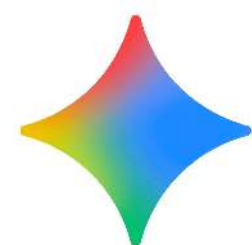
Claude

For the assistance in coding



ChatGPT

For discussion and emotional support



Gemini

For generating all comics in the presentation

Part 1: Introduction to LLM

Introduction: What is LLM



ChatGPT

hello

Thought for a couple of seconds >

Hi there! How can I help you today?



Question:

what is happening when you chat to ChatGPT/Claude/etc?

How are the text generated?

Introduction: What is LLM

Raw LLM (Foundation Model) is just a probability predictor



Nothing but a
conditional probability

$$P(x_t \mid x_1, x_2, \dots, x_{t-1})$$

Probability of next token given all previous tokens

Next word: sampled from the
learned probability distribution

What this is NOT

Understanding
No grounded semantics

Reasoning
No explicit logic

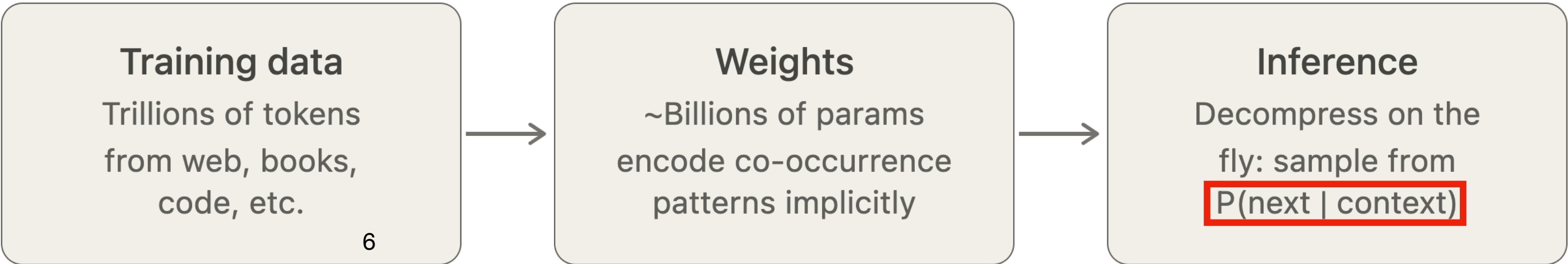
What this IS

Pattern matching
Over vast corpora

Compression
Statistical model

Analogy: extreme statistical compression

Training: learn the probability
distribution from data



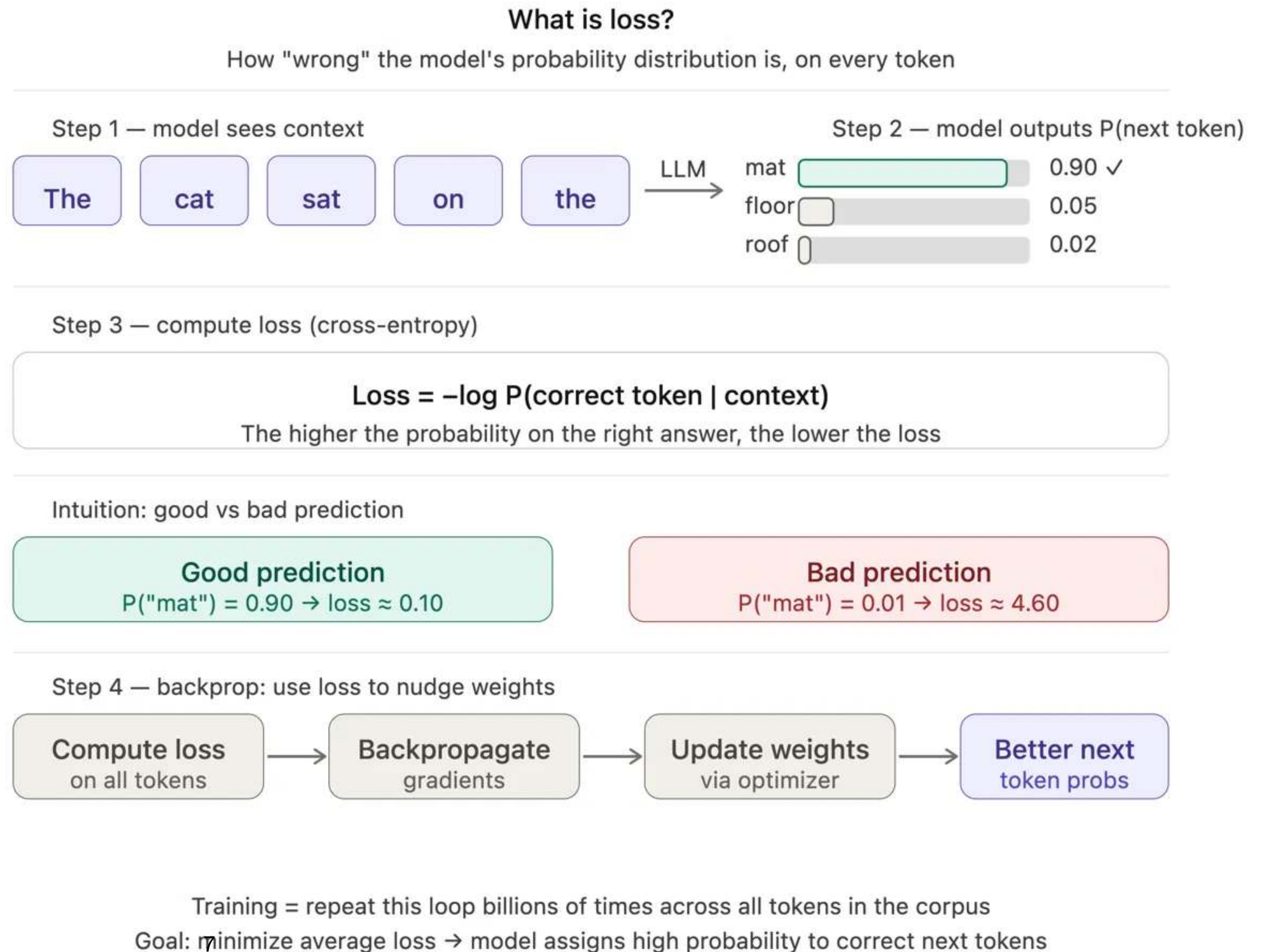
Introduction: What is LLM

Raw LLM (Foundation Model) is just a probability predictor

Nothing but a
conditional probability

**Next word: sampled from the
learned probability distribution**

**Training: learn the probability
distribution from data**

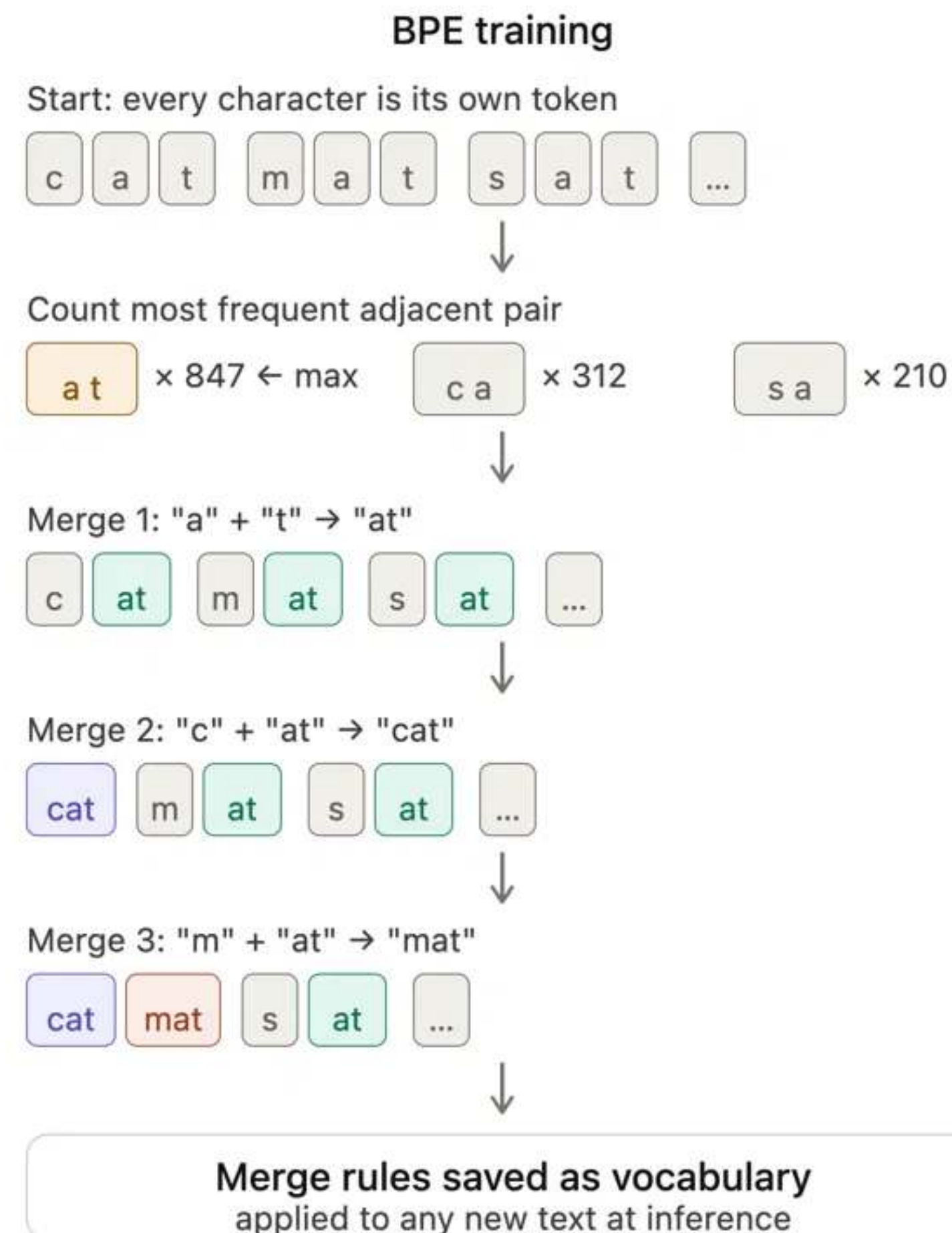


Introduction: What is LLM

Tokenization: Text -> Token

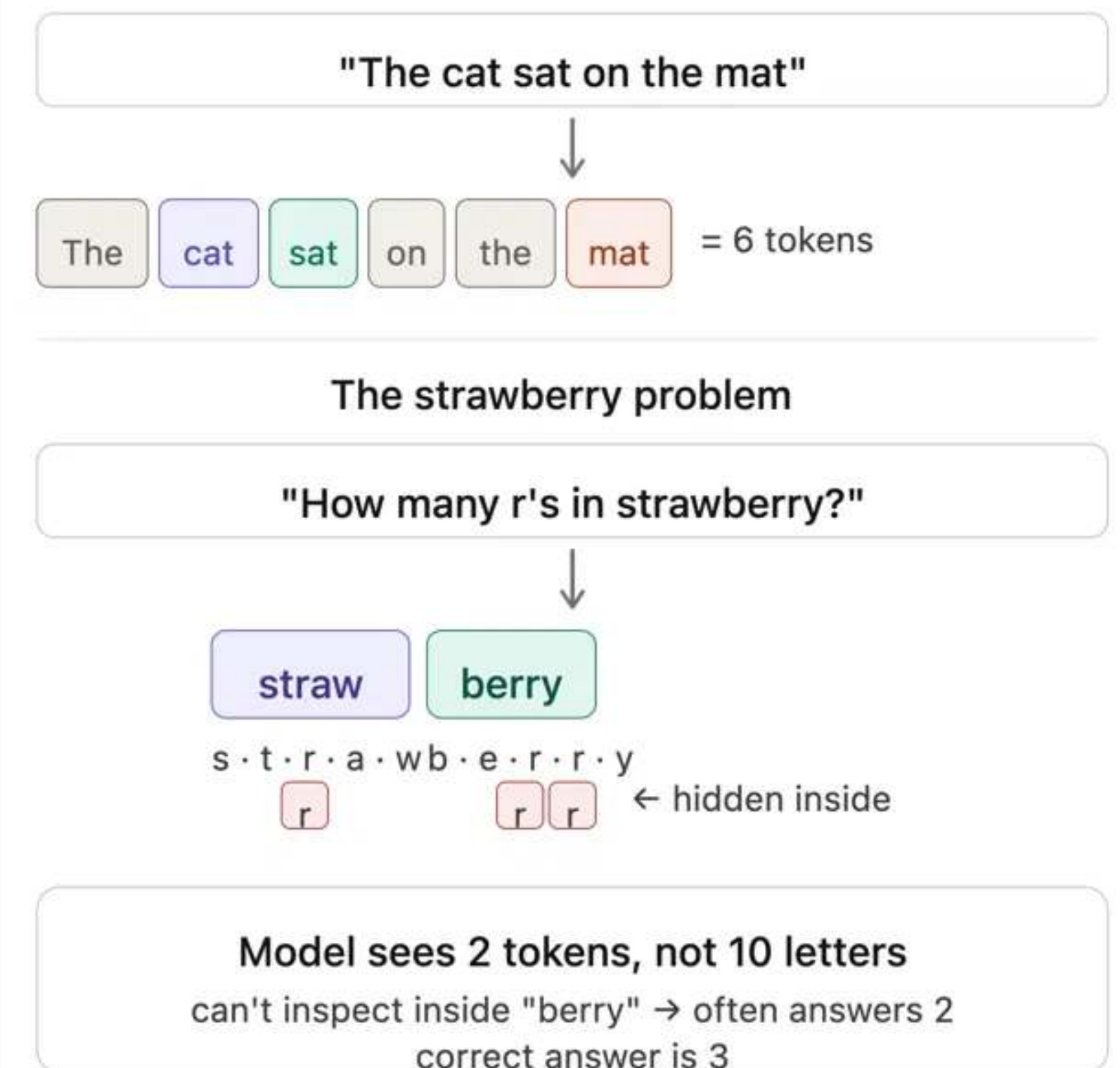
BPE: Byte Pair Encoding

- Not based on grammar, but purely **data-driven**
- Merge the **high-frequency combinations**
- Stop when the pool size is satisfactory (GPT-4 used 100k token)



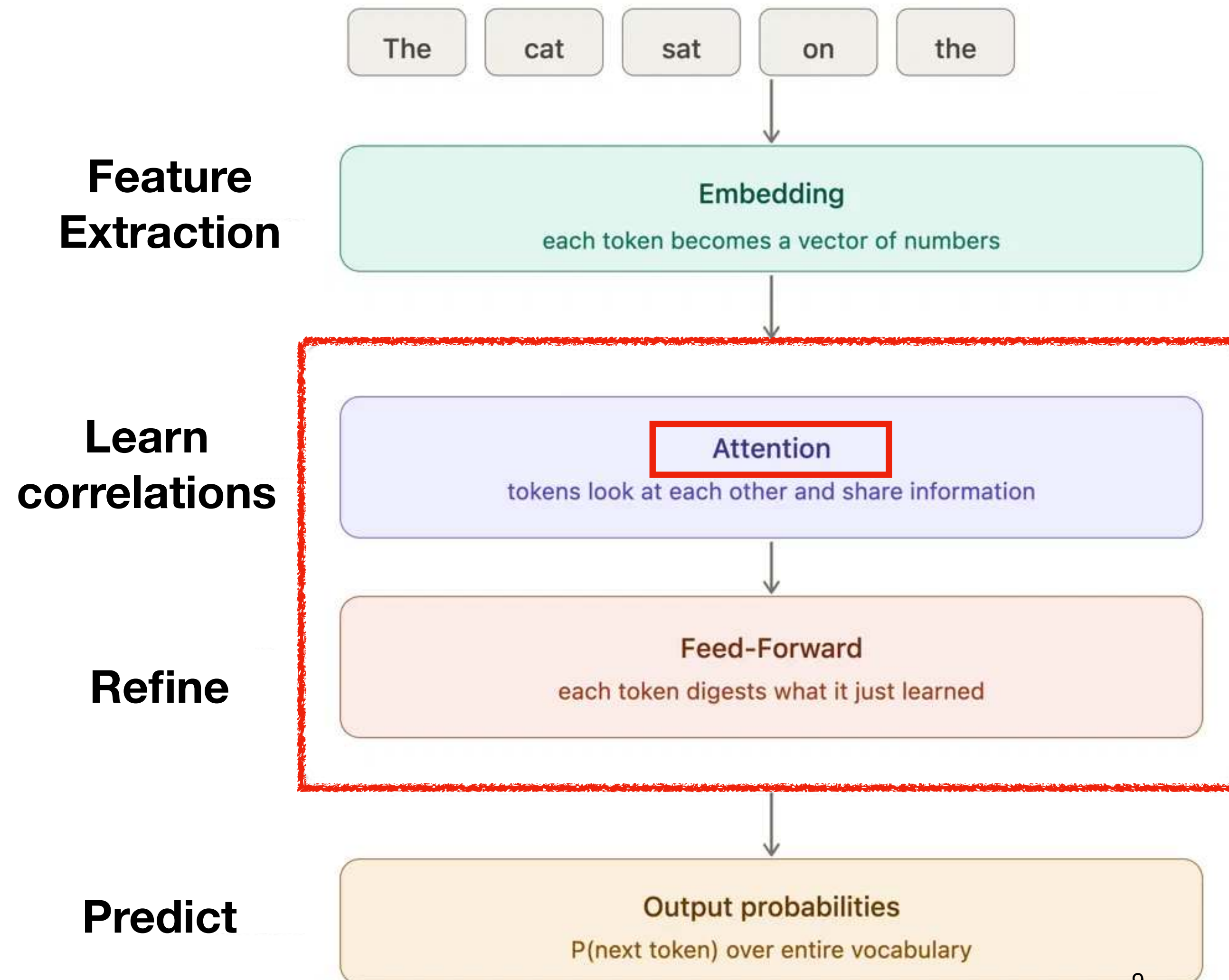
Question:
How about 1 token per character?
Or 1 token per word?

Encoding at inference



Introduction: What is LLM

Transformer: **Attention is all you need!**



Attention Is All You Need

Ashish Vaswani*
Google Brain
avaswani@google.com

Noam Shazeer*
Google Brain
noam@google.com

Niki Parmar*
Google Research
nikip@google.com

Jakob Uszkoreit*
Google Research
usz@google.com

Llion Jones*
Google Research
llion@google.com

Aidan N. Gomez* †
University of Toronto
aidan@cs.toronto.edu

Łukasz Kaiser*
Google Brain
lukaszkaizer@google.com

Illia Polosukhin* ‡
illia.polosukhin@gmail.com

Abstract

The dominant sequence transduction models are based on complex recurrent or convolutional neural networks that include an encoder and a decoder. The best performing models also connect the encoder and decoder through an attention mechanism. We propose a new simple network architecture, the Transformer, based solely on attention mechanisms, dispensing with recurrence and convolutions entirely. Experiments on two machine translation tasks show these models to be superior in quality while being more parallelizable and requiring significantly less time to train. Our model achieves 28.4 BLEU on the WMT 2014 English-to-German translation task, improving over the existing best results, including ensembles, by over 2 BLEU. On the WMT 2014 English-to-French translation task, our model establishes a new single-model state-of-the-art BLEU score of 41.8 after training for 3.5 days on eight GPUs, a small fraction of the training costs of the best models from the literature. We show that the Transformer generalizes well to other tasks by applying it successfully to English constituency parsing both with large and limited training data.

Introduction: What is LLM

Transformer: Attention is all you need!

For each token:

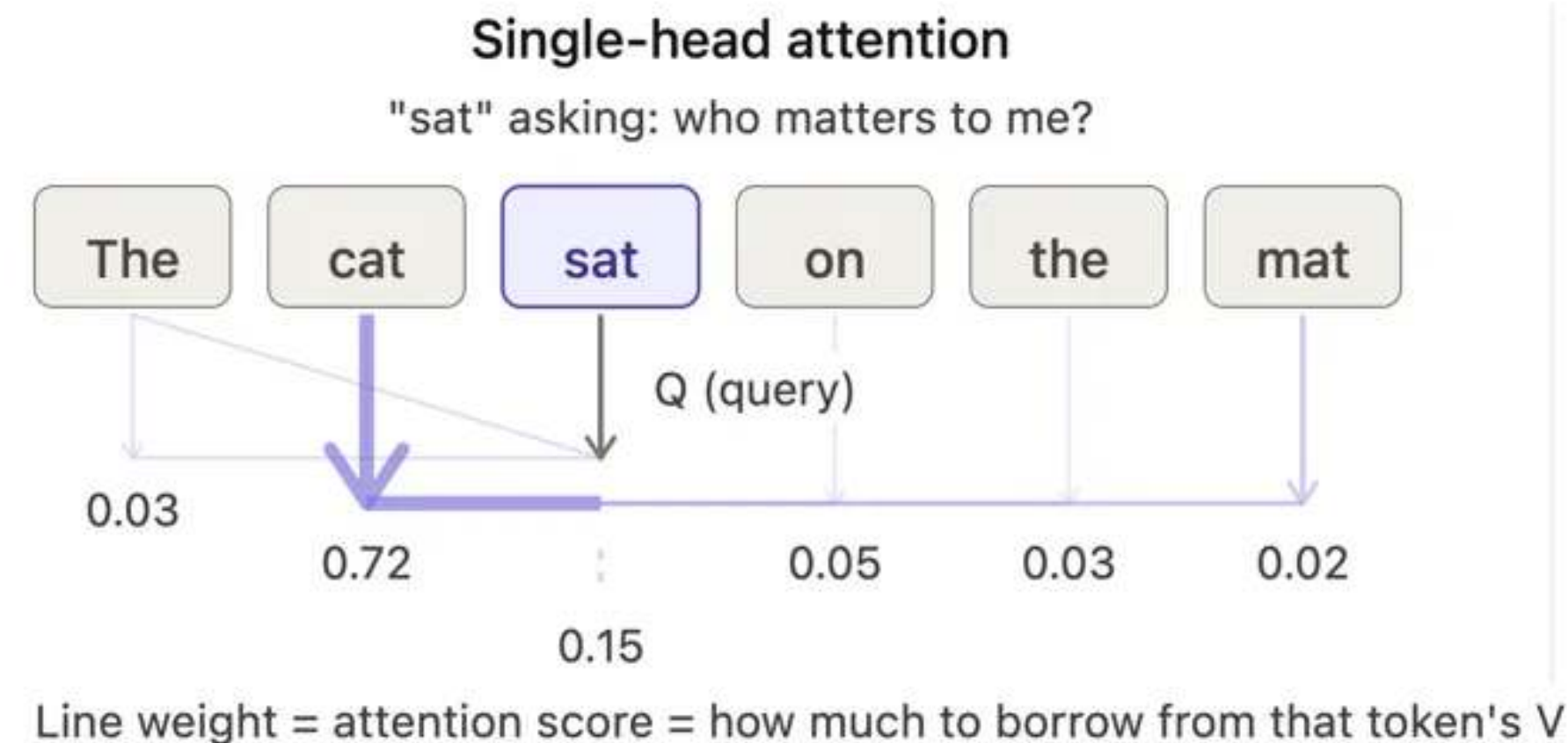
- Q, K, V are **vectors**

For a sentence (array):

- Q, K, V are **matrix**

Self Attention:

- Tokens “query” each other

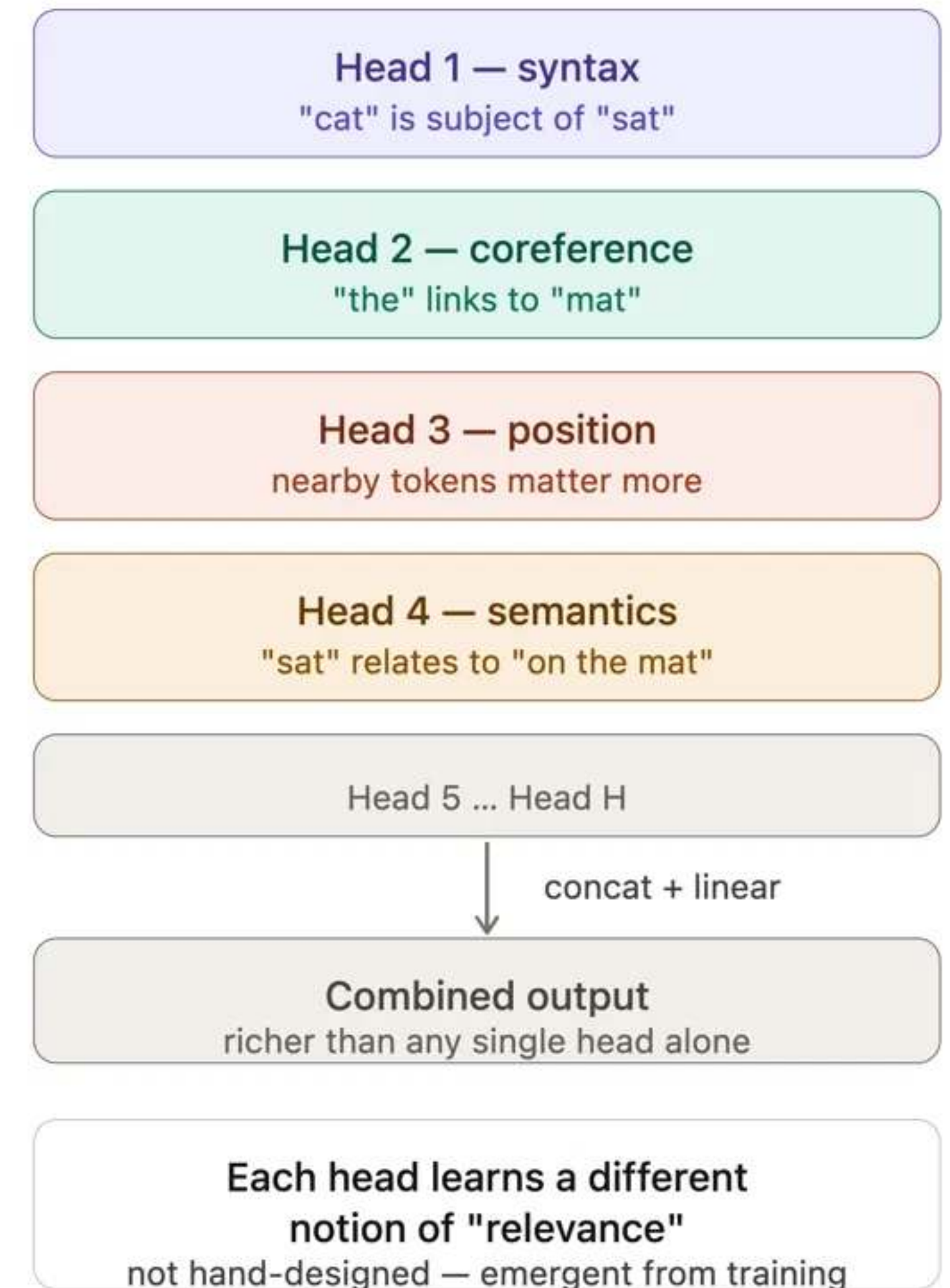


$$\text{softmax}(Q \cdot K^T / \sqrt{d_k}) \cdot V$$

$Q \cdot K^T \rightarrow$ compatibility score between query and every key
 $\div \sqrt{d_k} \rightarrow$ scale down to stabilise gradients
 $\text{softmax} \rightarrow$ turn scores into weights that sum to 1
 $\cdot V \rightarrow$ weighted mix of value vectors

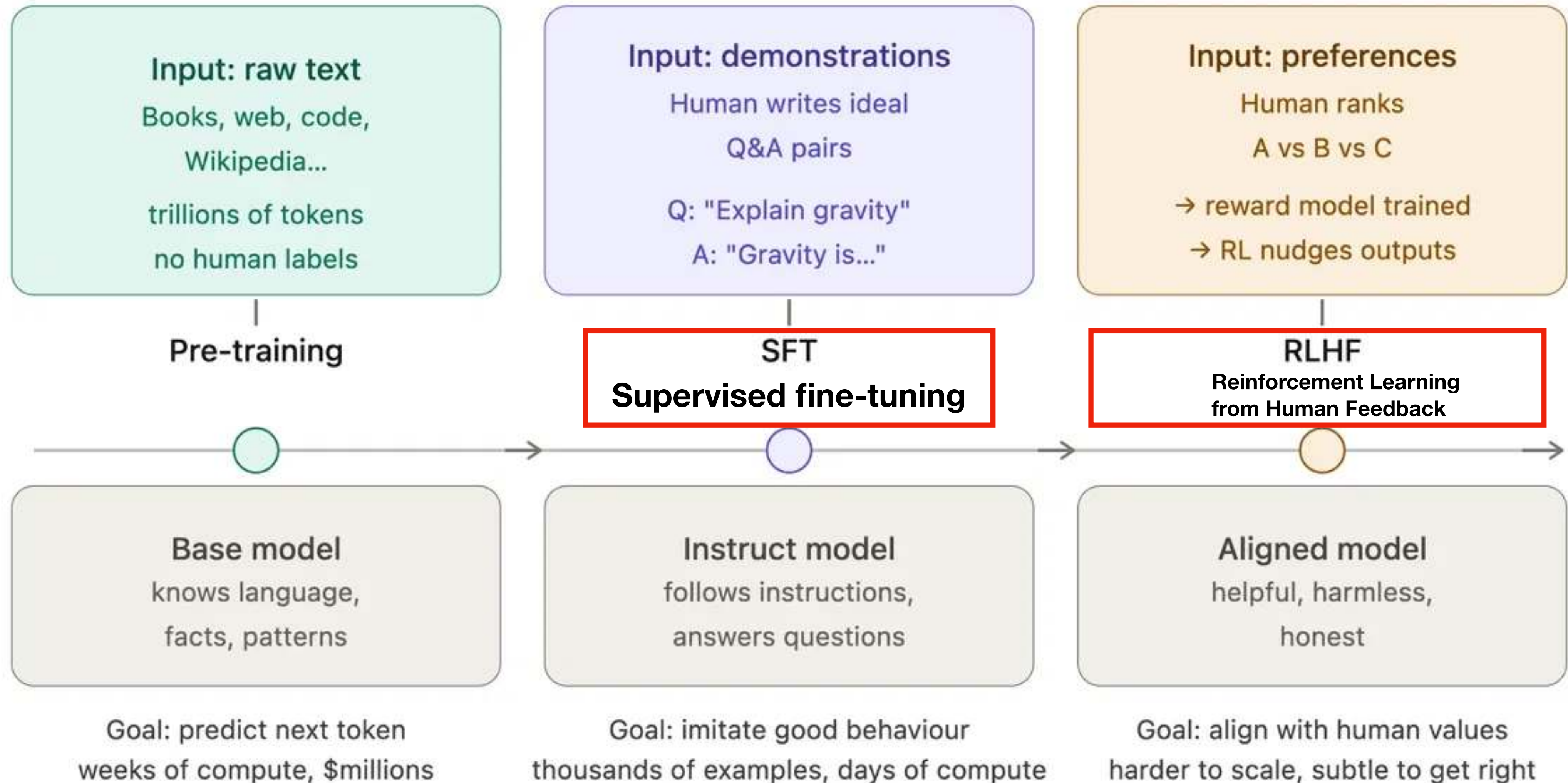
Multi-head attention

run H attention heads in parallel



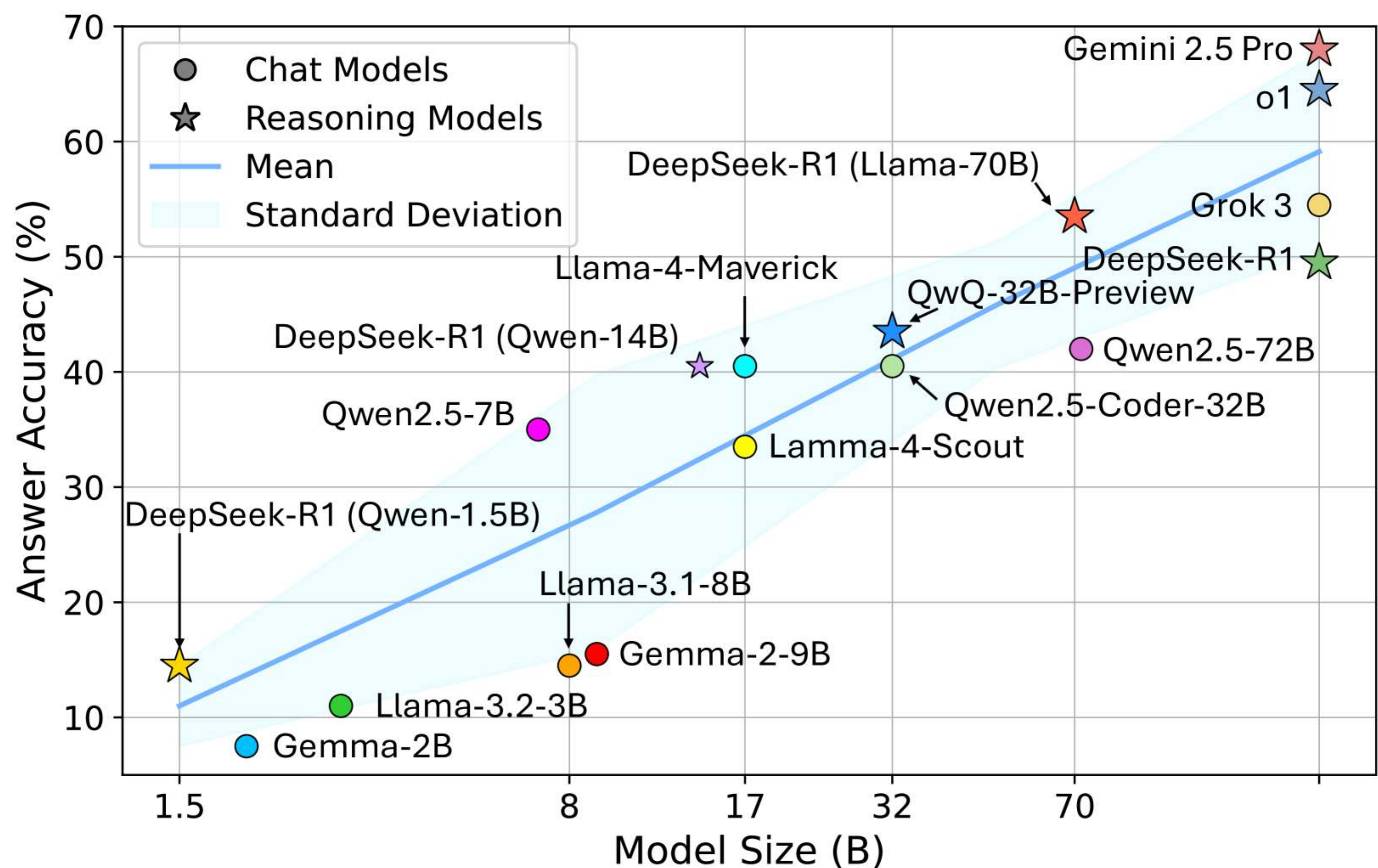
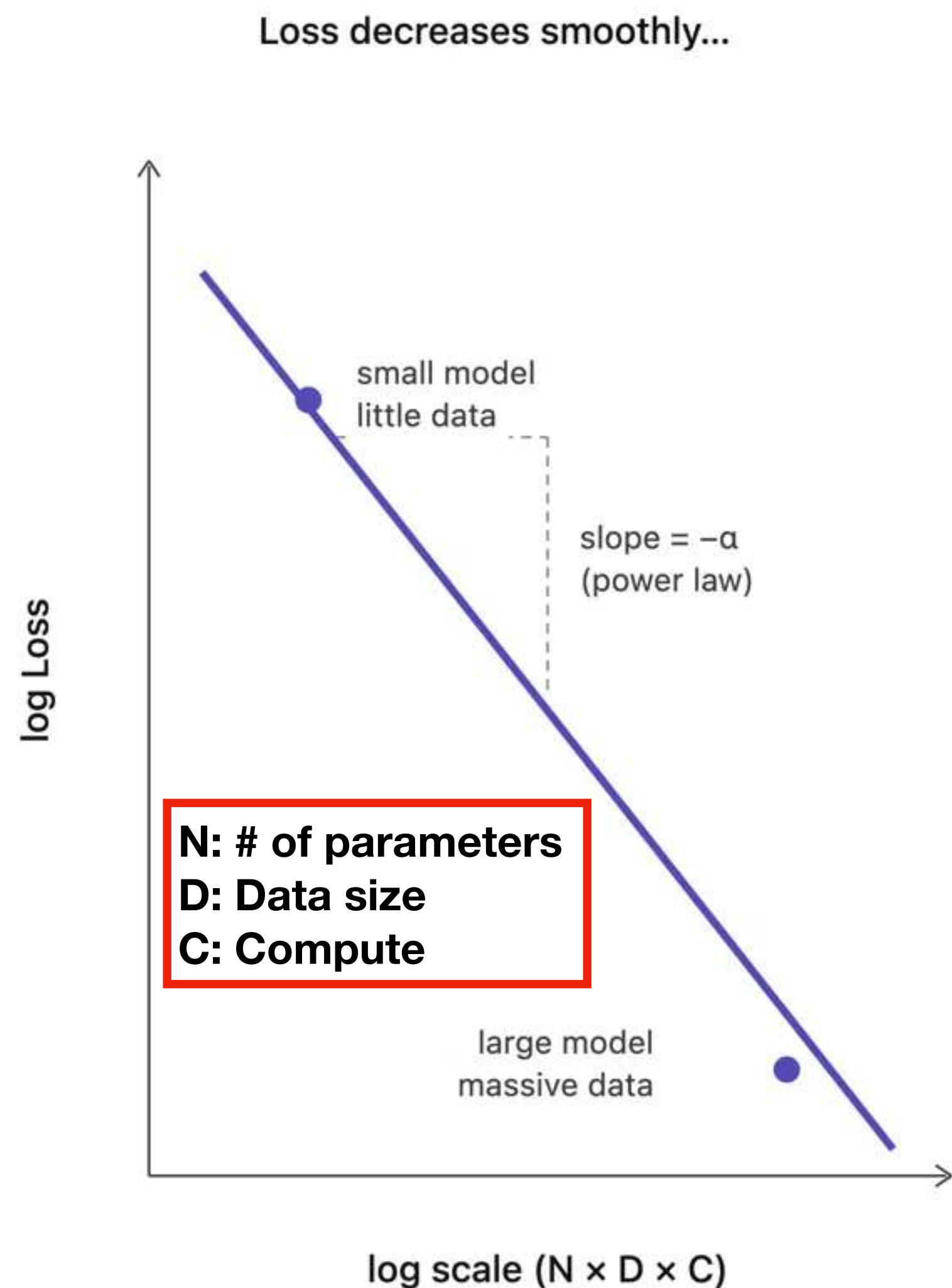
Lifecycle of a Raw LLM

Pre-training -> Post training



Lifecycle of a Raw LLM

The Scaling Law



From Raw LLM to Chat Model

Previous question: what we are actually chatting with?

Raw LLM

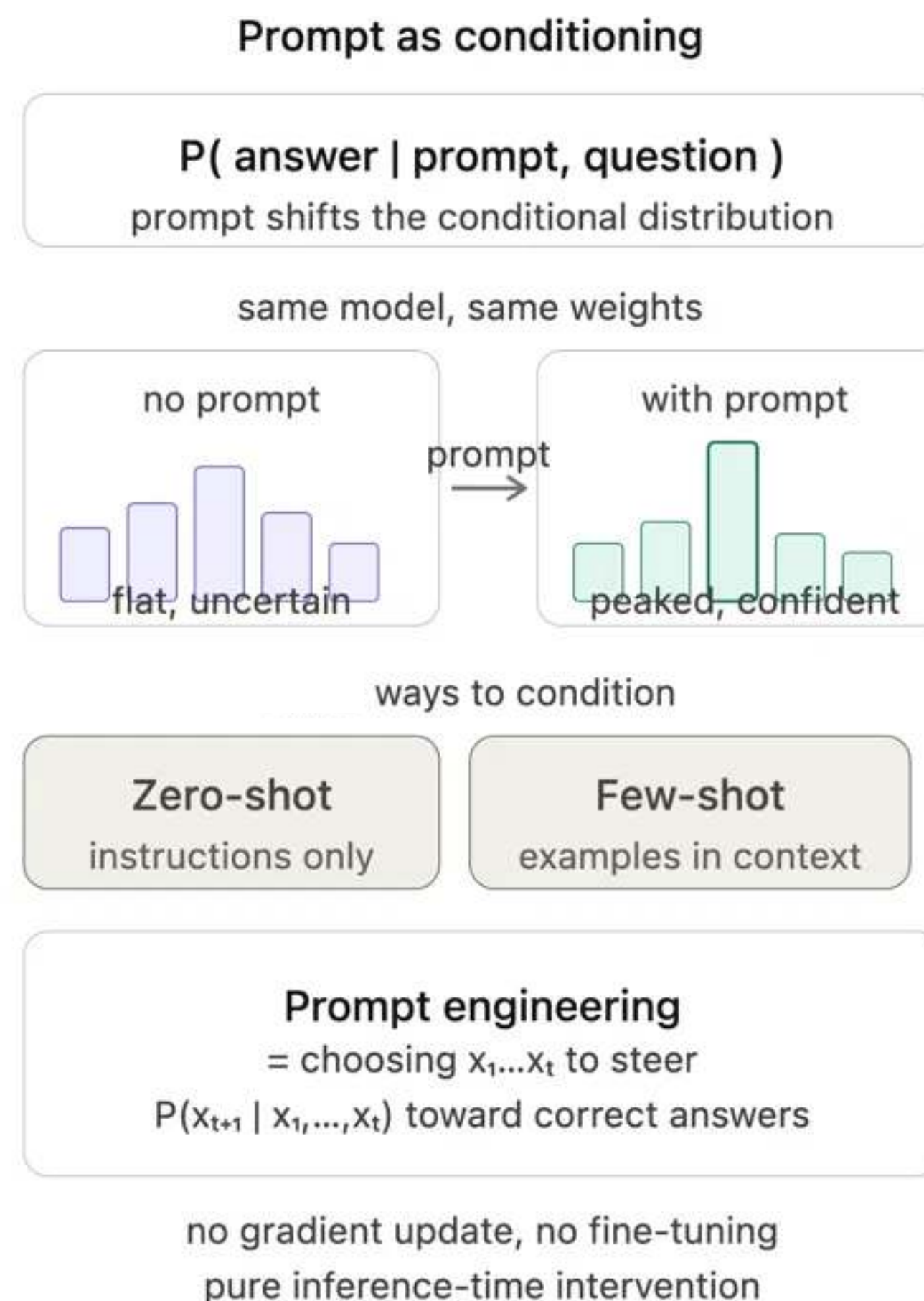
+

- **Prompt Engineering**
- **Knowledge injection**
- **Test-time compute (Reasoning, etc)**
- **Tool calling**
- **Memory control**
- **And more...**

From Raw LLM to Chat Model

Prompt Engineering

- **Better condition** provided to the probability predictor
- **Few-shot**: provide some examples
- **Zero-shot**: just some high-level principle!



Two examples

① Zero-shot CoT

Kojima et al., 2022 · arXiv:2205.11916

5 Magical Words

without CoT

"Q: Roger has 5 balls.
He buys 2 more cans of 3.
How many balls?"

"A: 11" ×

MultiArith: 17.7%

+ "Let's think step by step"

"Q: Roger has 5 balls.
... Let's think step by step."
"He starts with 5,
buys 2×3=6, total 11" ✓

MultiArith: 78.7%

+61% accuracy · same model · 5 words added

② Few-shot prompting

Brown et al., 2020 · GPT-3 paper · arXiv:2005.14165

zero-shot

"Translate to French:"

"cheese"

"fromage" ✓

(simple task: ok)

(complex task: fails)

few-shot (examples in prompt)

"sea otter → loutre de mer"

"peppermint → menthe"

"plush giraffe → girafe"

"cheese →

fromage" ✓✓

64.3% → 71.2% · no fine-tuning · examples only

From Raw LLM to Chat Model

Prompt Engineering



prompts.chat

The world's largest open-source prompt library for AI

Works with ChatGPT, Claude, Gemini, Llama, Mistral, and more

formerly known as Awesome ChatGPT Prompts



[Browse Prompts](#) • [Read the Book](#) • [View on GitHub](#) • [Self-Host](#)

🏆 Featured in [Forbes](#) • 🎓 Referenced by [Harvard](#), [Columbia](#) • 📄 [40+ academic citations](#) • ❤️ [Most liked dataset](#) on Hugging Face

★ 143k+ GitHub stars • 🏆 [GitHub Staff Pick](#) • 🚀 First prompt library (Dec 2022)

Loved by AI pioneers:

[Greg Brockman](#) (OpenAI Co-Founder) • [Wojciech Zaremba](#) (OpenAI Co-Founder) • [Clement Delangue](#) (Hugging Face CEO) •

[Thomas Dohmke](#) (Former GitHub CEO)

<https://github.com/f/prompts.chat>

From Raw LLM to Chat Model

Test-time compute

can you check the whether today for me?

Thought for 8s >

Currently 70° · Cloudy

Coronado, CA, United States



Today

70° 64°

After a cloudy start, sun returns



Thinking

- **Determining user location for weather query**
The user likely meant "weather" instead of "whether."
They're in Coronado, California, based on available info, but I should confirm their exact location to ensure accuracy. I'll check this before retrieving the weather details to provide an accurate response.
- **Generating weather forecast response with UI**
I need to call for the user's location first, likely "Coronado, CA," then retrieve the weather forecast for that area. I'll display today's weather using a forecast widget.
- **Confirming location and weather information**
Today
- **Getting weather for Coronado, CA**
Now, I'll fetch the weather for Coronado, California. I'll set the location to Coronado, CA and use the default settings for start and duration. With that, I can provide today's weather forecast. Let's get those details for the user!

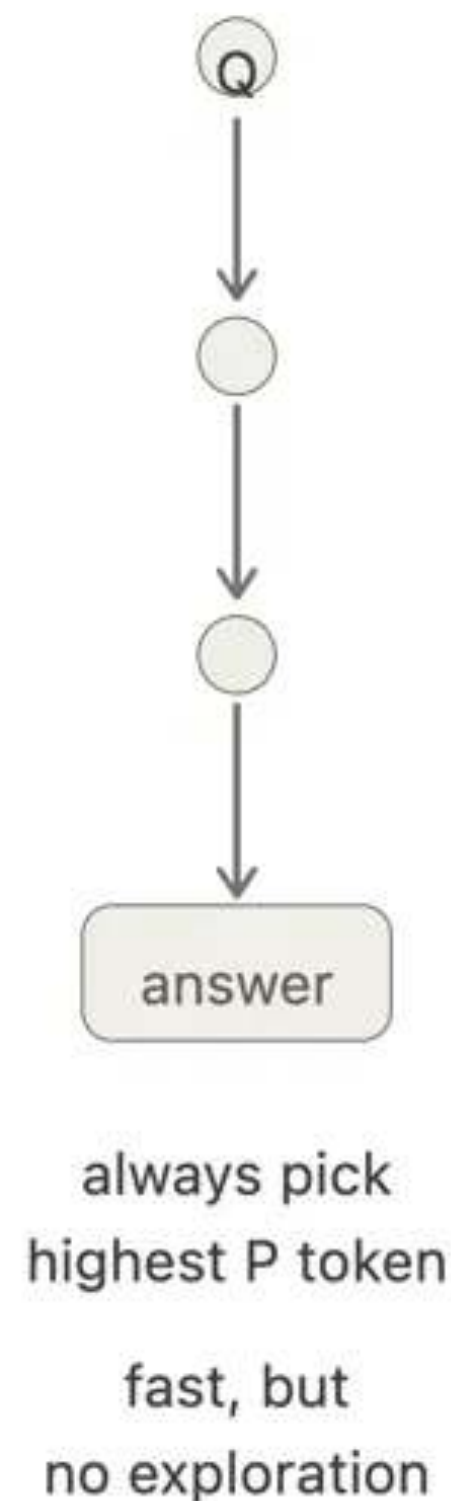
☑ Thought for 8s
Done

From Raw LLM to Chat Model

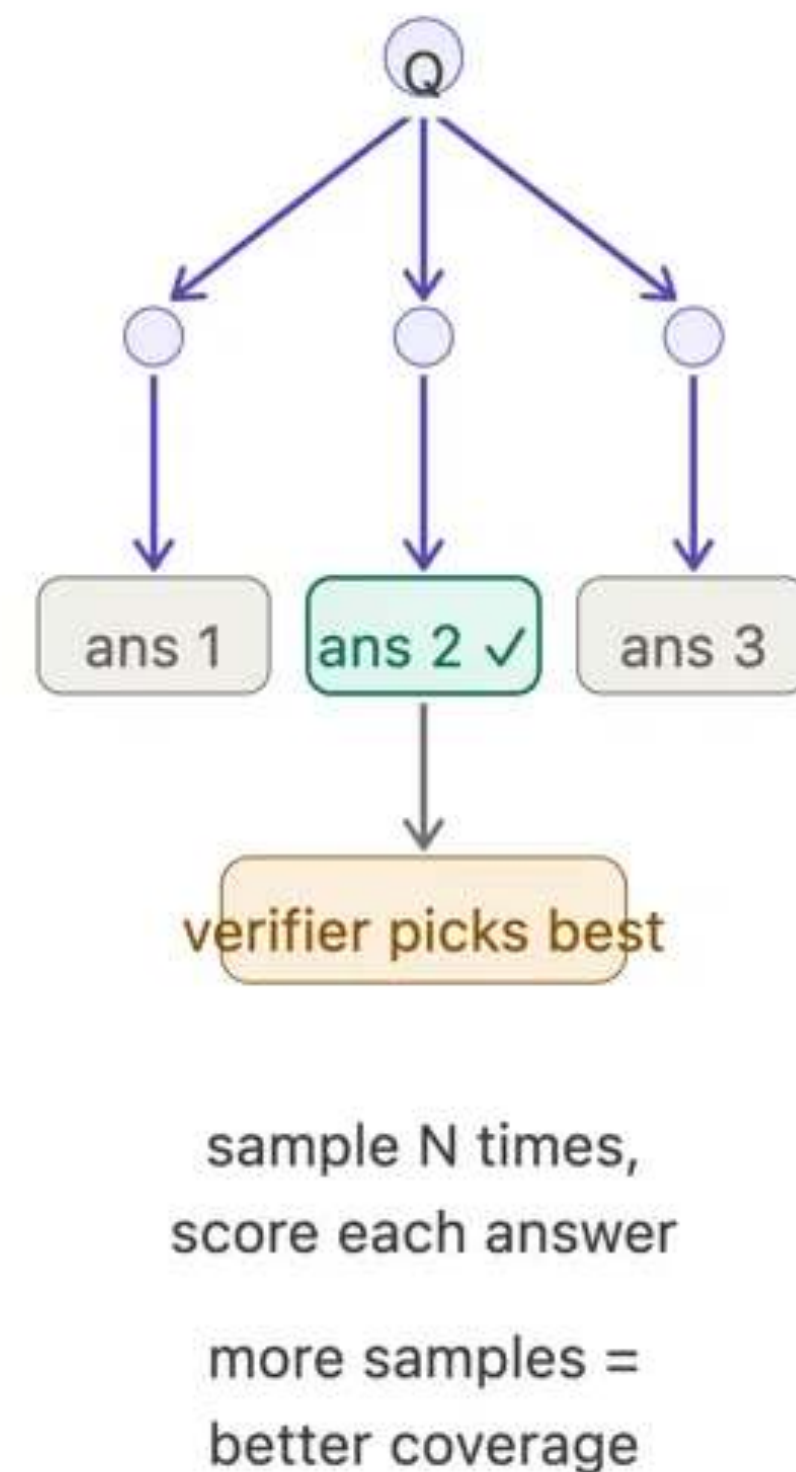
Test-time compute

Raw LLM = **Probability Distribution**, Test-time compute = **smarter sampling**

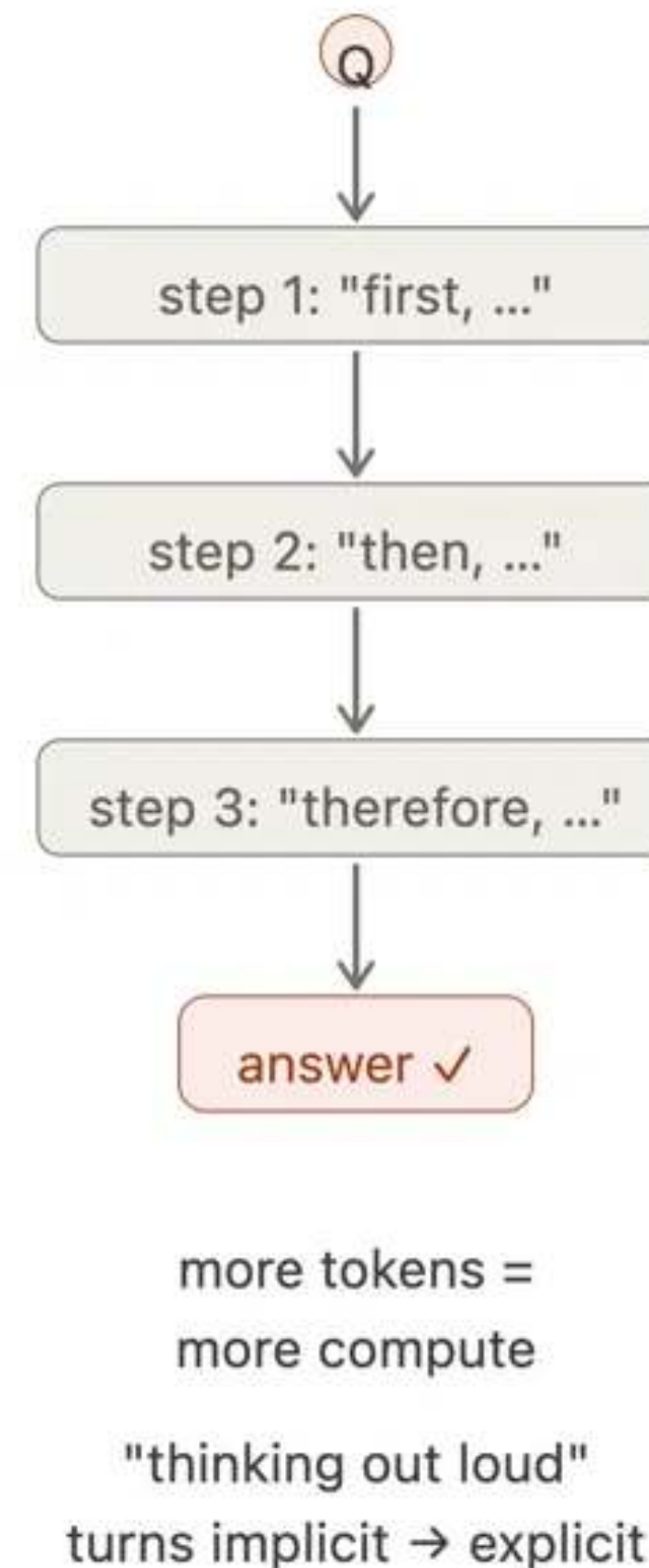
Greedy



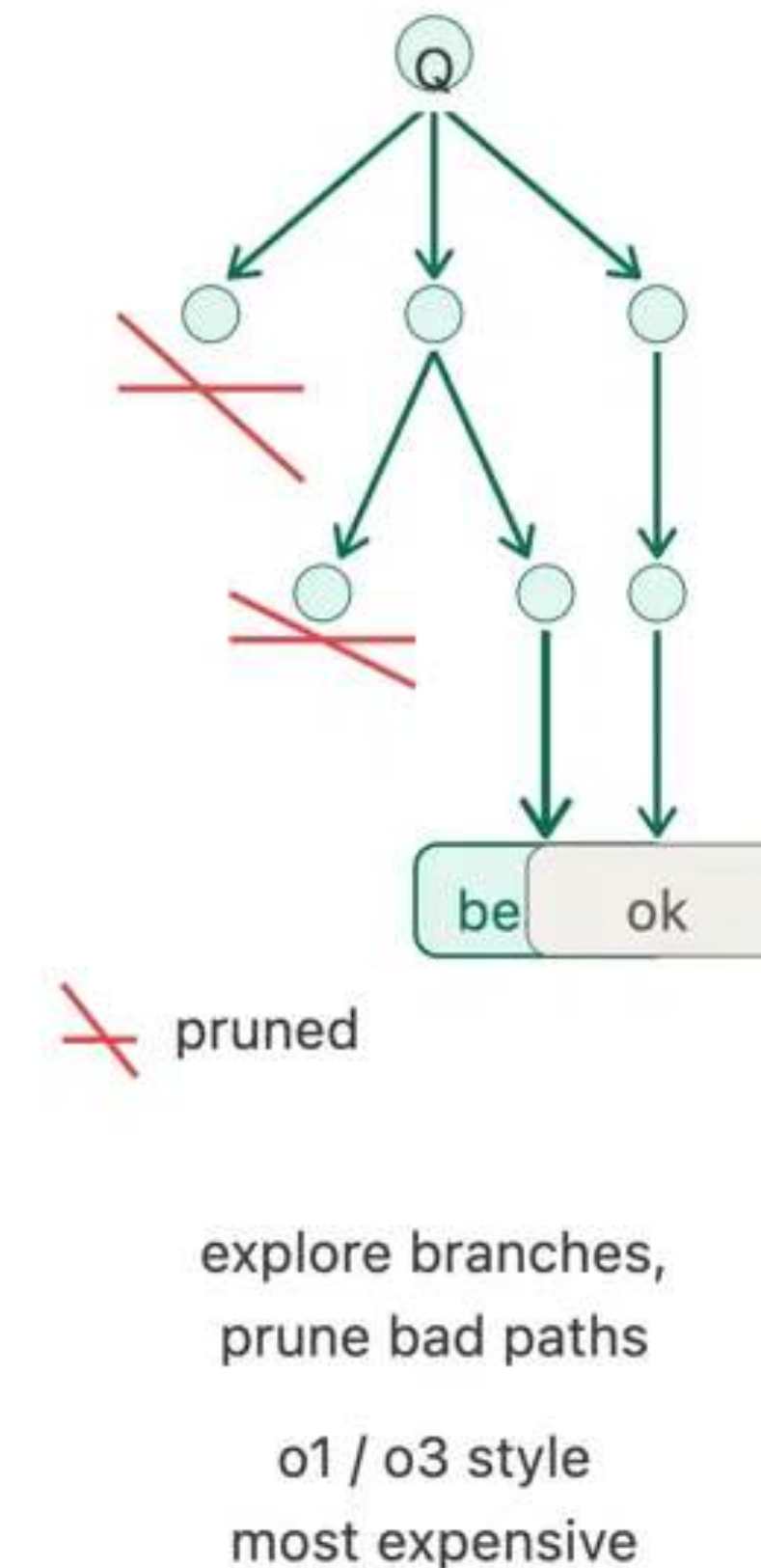
Best-of-N



Chain-of-Thought



Tree Search



From Raw LLM to Chat Model

Tool Calling

which tools did you use? i'm not sure whether to trust your answer!

Thought for 10s >

Which tools I used

I rely on a couple of built-in utilities:

1. **Location helper** – a private tool that just tells me your coarse location (e.g., Coronado, California) and the current date/time.
2. **Weather lookup** – a tool that queries a weather provider for the latest observations and 7-day forecast for that location.



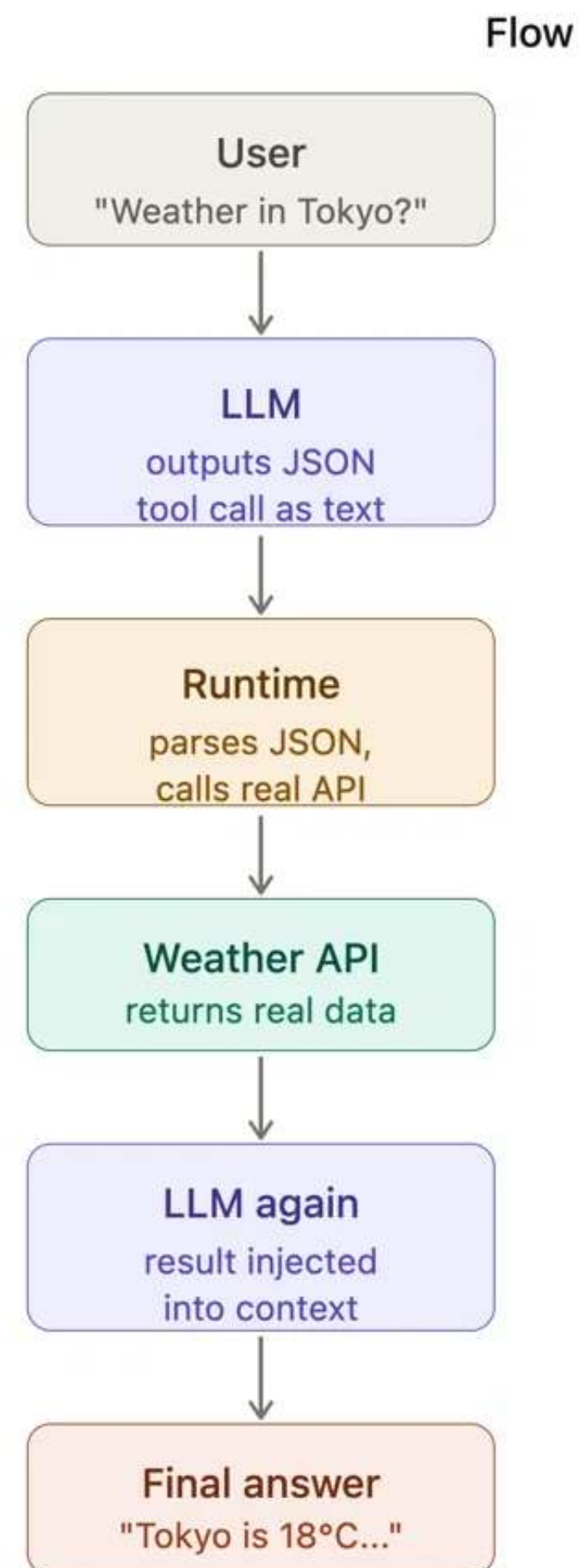
So smart!

But how did it happen?

I thought LLM can only produce some texts!

From Raw LLM to Chat Model

Tool Calling



Tool calling =

Pre-trained format (message -> JSON)

+

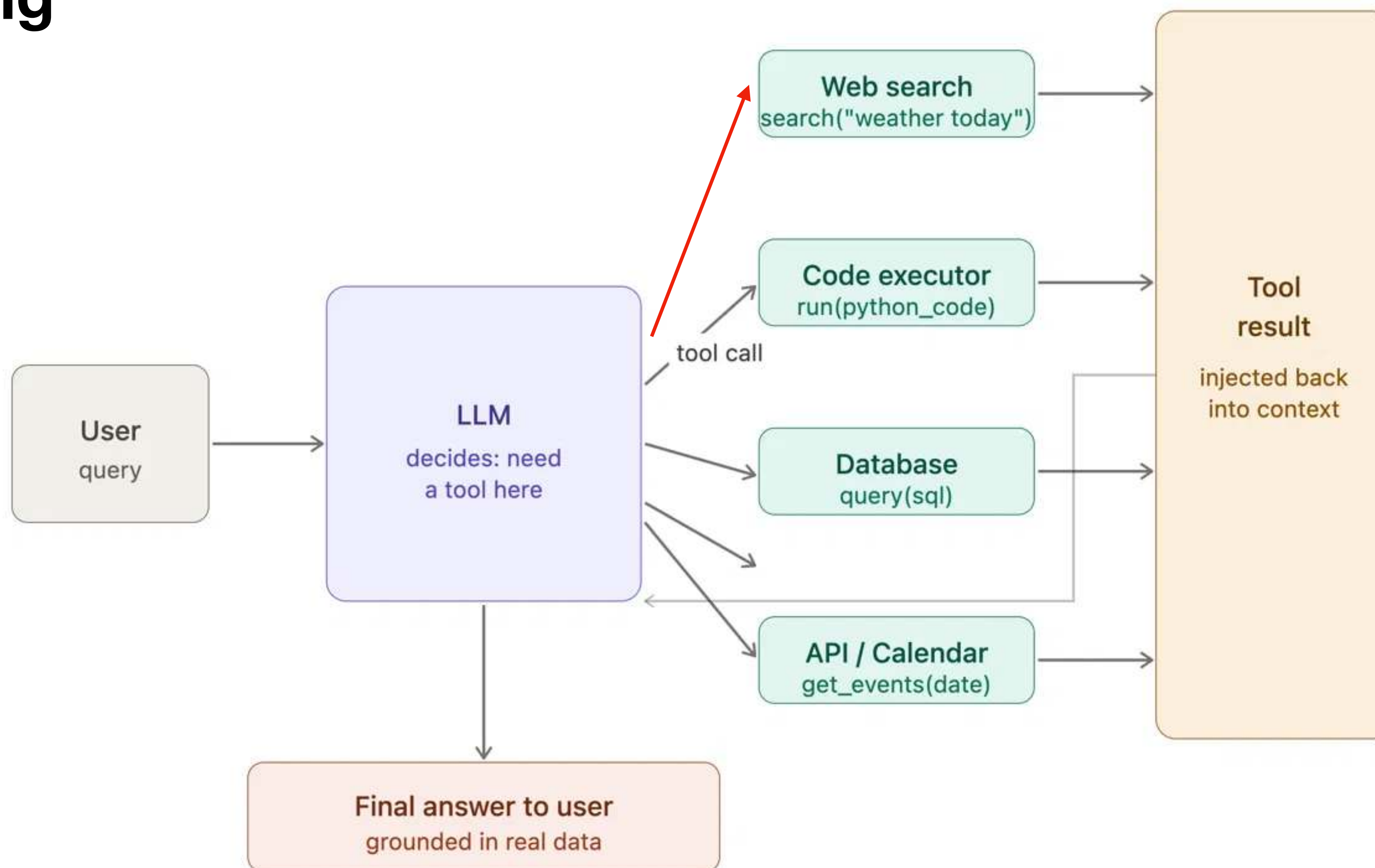
API (JSON -> real function call)

+

Context aggregation (result -> answer)

From Raw LLM to Chat Model

Tool Calling



LLM generates a structured "tool call" as text → runtime executes it → result appended to context → LLM continues

From Raw LLM to Chat Model

Memory Control



Partly solved by larger context window size (>=1M token)

How to Evaluate the LLMs?

Benchmarking is critical

Published as a conference paper at ICLR 2024

SWE-BENCH: CAN LANGUAGE MODELS RESOLVE REAL-WORLD GITHUB ISSUES?

Carlos E. Jimenez^{* 1,2} John Yang^{* 1,2} Alexander Wettig^{1,2}
Shunyu Yao^{1,2} Kexin Pei³ Ofir Press^{1,2} Karthik Narasimhan^{1,2}

¹Princeton University ²Princeton Language and Intelligence ³University of Chicago

ABSTRACT

Language models have outpaced our ability to evaluate them effectively, but for their future development it is essential to study the frontier of their capabilities. We find real-world software engineering to be a rich, sustainable, and challenging testbed for evaluating the next generation of language models. To this end, we introduce SWE-bench, an evaluation framework consisting of 2,294 software engineering problems drawn from real GitHub issues and corresponding pull requests across 12 popular Python repositories. Given a codebase along with a description of an issue to be resolved, a language model is tasked with editing the codebase to address the issue. Resolving issues in SWE-bench frequently requires understanding and coordinating changes across multiple functions, classes, and even files simultaneously, calling for models to interact with execution environments, process extremely long contexts and perform complex reasoning that goes far beyond traditional code generation tasks. Our evaluations show that both state-of-the-art proprietary models and our fine-tuned model SWE-Llama can resolve only the simplest issues. The best-performing model, Claude 2, is able to solve a mere 1.96% of the issues. Advances on SWE-bench represent steps towards LMs that are more practical, intelligent, and autonomous.

SWE-bench

Official Leaderboards

SWE-agent-LM-32B, the open-weight SotA (4/2025) on Verified, trained on synthetic data generated by SWE-smith.
[More in the paper!](#)

VerifiedMultilingualLiteFullMultimodal

Verified is a human-filtered subset of 500 instances. We use [mini-SWE-agent](#) to evaluate all models with the same harness ([details](#)).

Compare results

Agent: mini-SWE-agent v2Models: All models

☐	Model	% Resolved	Avg. \$	Trajs	Org	Date
☐	NEW Claude 4.5 Opus (high reasoning)	76.80	\$0.75	🔗	AI	2026-02-17
☐	NEW Gemini 3 Flash (high reasoning)	75.80	\$0.36	🔗	🌟	2026-02-17
☐	NEW MiniMax M2.5 (high reasoning)	75.80	\$0.07	🔗	📺	2026-02-17
☐	NEW Claude Opus 4.6	75.60	\$0.55	🔗	AI	2026-02-17
☐	NEW GPT-5-2 Codex	72.80	\$0.45	🔗	🌀	2026-02-19

How to Evaluate the LLMs?

Benchmarking is critical

Humanity's Last Exam

Organizing Team

Long Phan^{*1}, Alice Gatti^{*1}, Ziwen Han^{*2}, Nathaniel Li^{*1},

Josephina Hu², Hugh Zhang[†], Chen Bo Calvin Zhang², Mohamed Shaaban², John Ling², Sean Shi², Michael Choi², Anish Agrawal², Arnav Chopra², Adam Khoja¹, Ryan Kim[†], Richard Ren¹, Jason Hausenloy¹, Oliver Zhang¹, Mantas Mazeika¹,

Summer Yue^{**2}, Alexandr Wang^{**2}, Dan Hendrycks^{**1}

¹ Center for AI Safety, ² Scale AI

Abstract

Benchmarks are important tools for tracking the rapid advancements in large language model (LLM) capabilities. However, benchmarks are not keeping pace in difficulty: LLMs now achieve over 90% accuracy on popular benchmarks like MMLU, limiting informed measurement of state-of-the-art LLM capabilities. In response, we introduce HUMANITY'S LAST EXAM (HLE), a multi-modal benchmark at the frontier of human knowledge, designed to be the final closed-ended academic benchmark of its kind with broad subject coverage. HLE consists of 2,500 questions across dozens of subjects, including mathematics, humanities, and the natural sciences. HLE is developed globally by subject-matter experts and consists of multiple-choice and short-answer questions suitable for automated grading. Each question has a known solution that is unambiguous and easily verifiable, but cannot be quickly answered via internet retrieval. State-of-the-art LLMs demonstrate low accuracy and calibration on HLE, highlighting a significant gap between current LLM capabilities and the expert human frontier on closed-ended academic questions. To inform research and policymaking upon a clear understanding of model capabilities, we publicly release HLE at <https://lastexam.ai>.

2500 PhD-level Q&A



Designed for O(years) of development

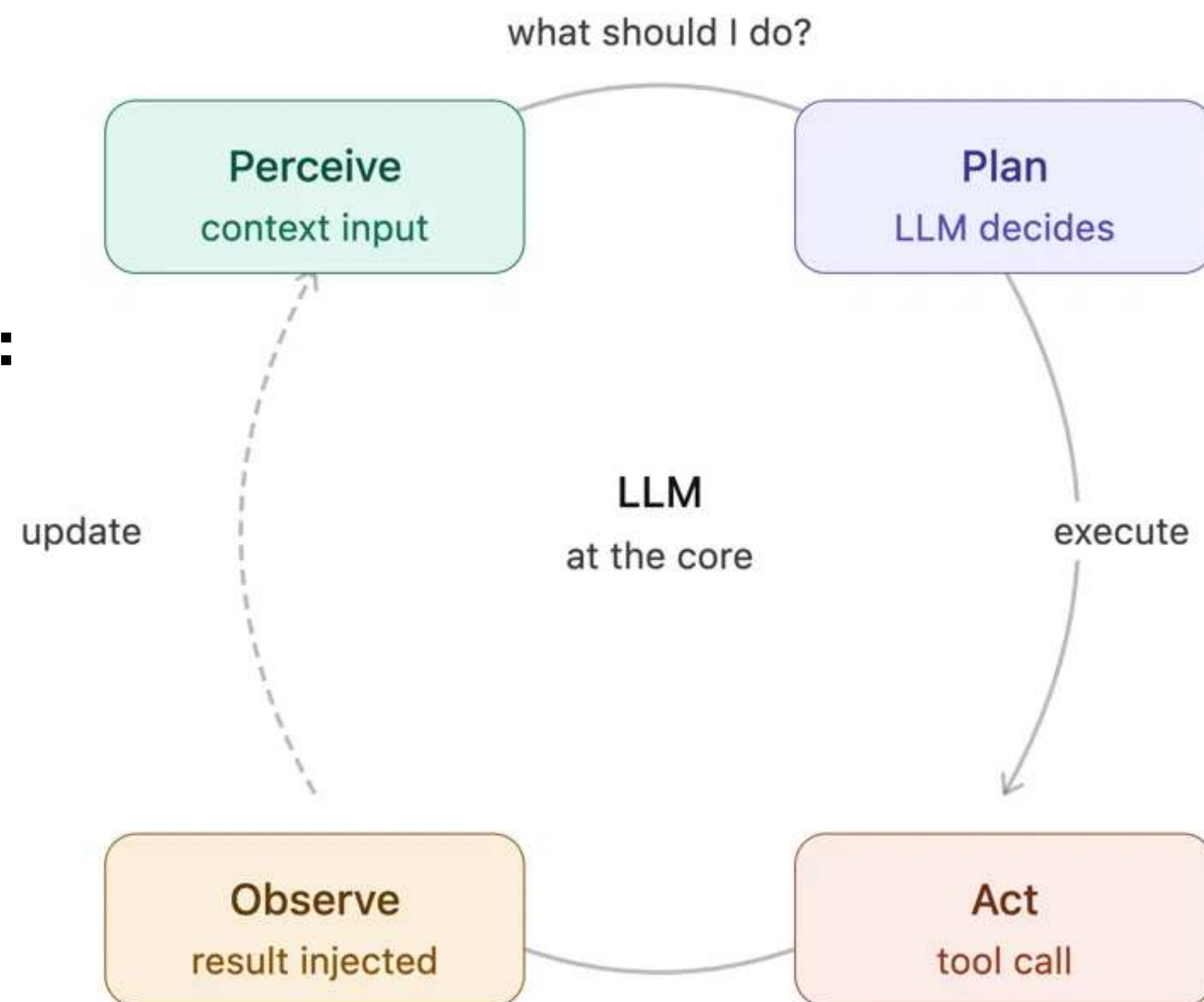
But it's saturated very quickly

New benchmarks are always needed

From Chat Model to Agent

Definition of LLM Agent

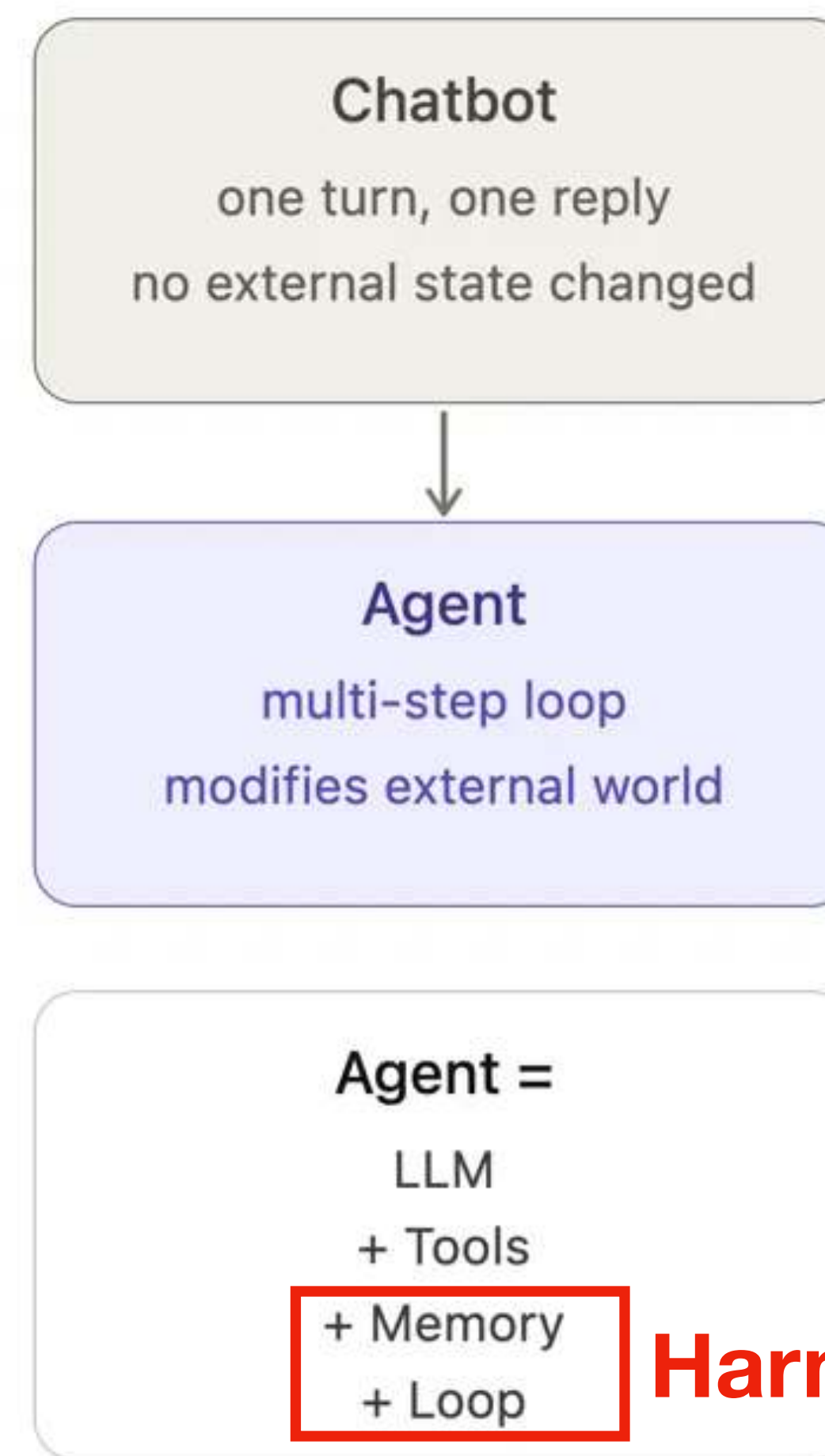
The agent loop



Beyond the chat model:

- Multi-step loop
- Long-term memory

vs chatbot

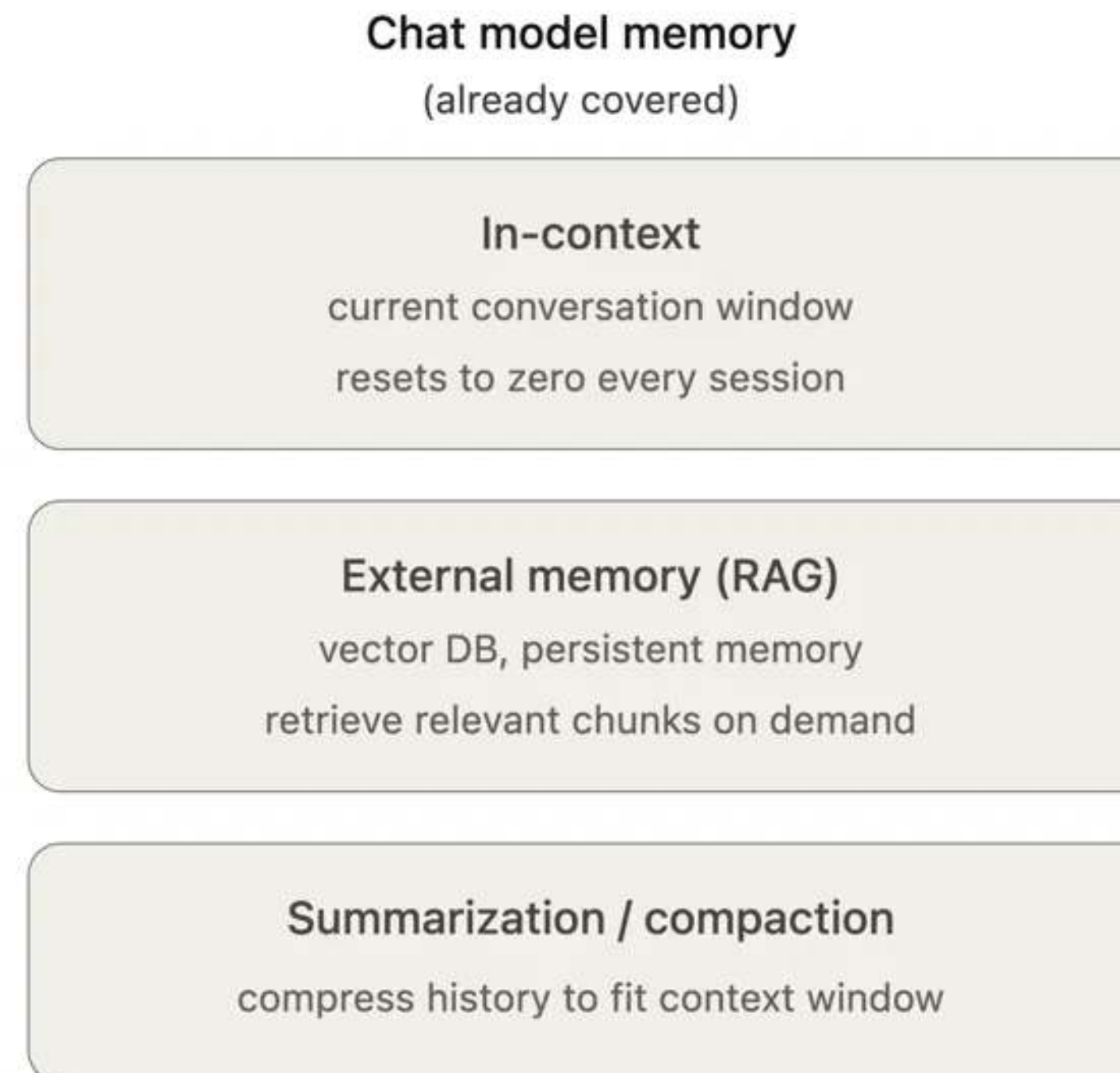


From Chat Model to Agent

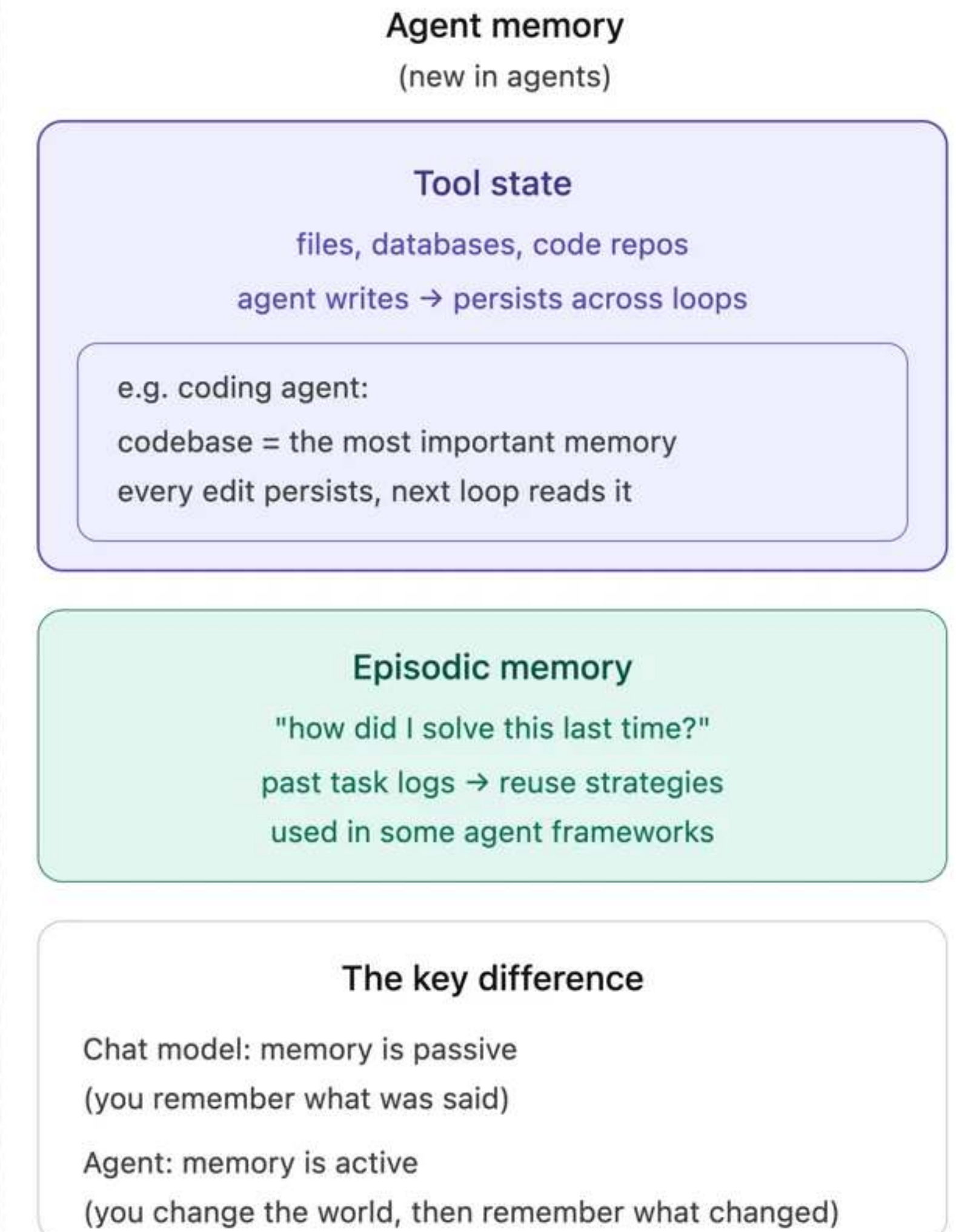
Memory of LLM Agents

Beyond the chat model:

- Multi-step loop: **tool state**
itself acts as a long-term
memory
grep instead of RAG
- Long-term memory:
transferrable between tasks

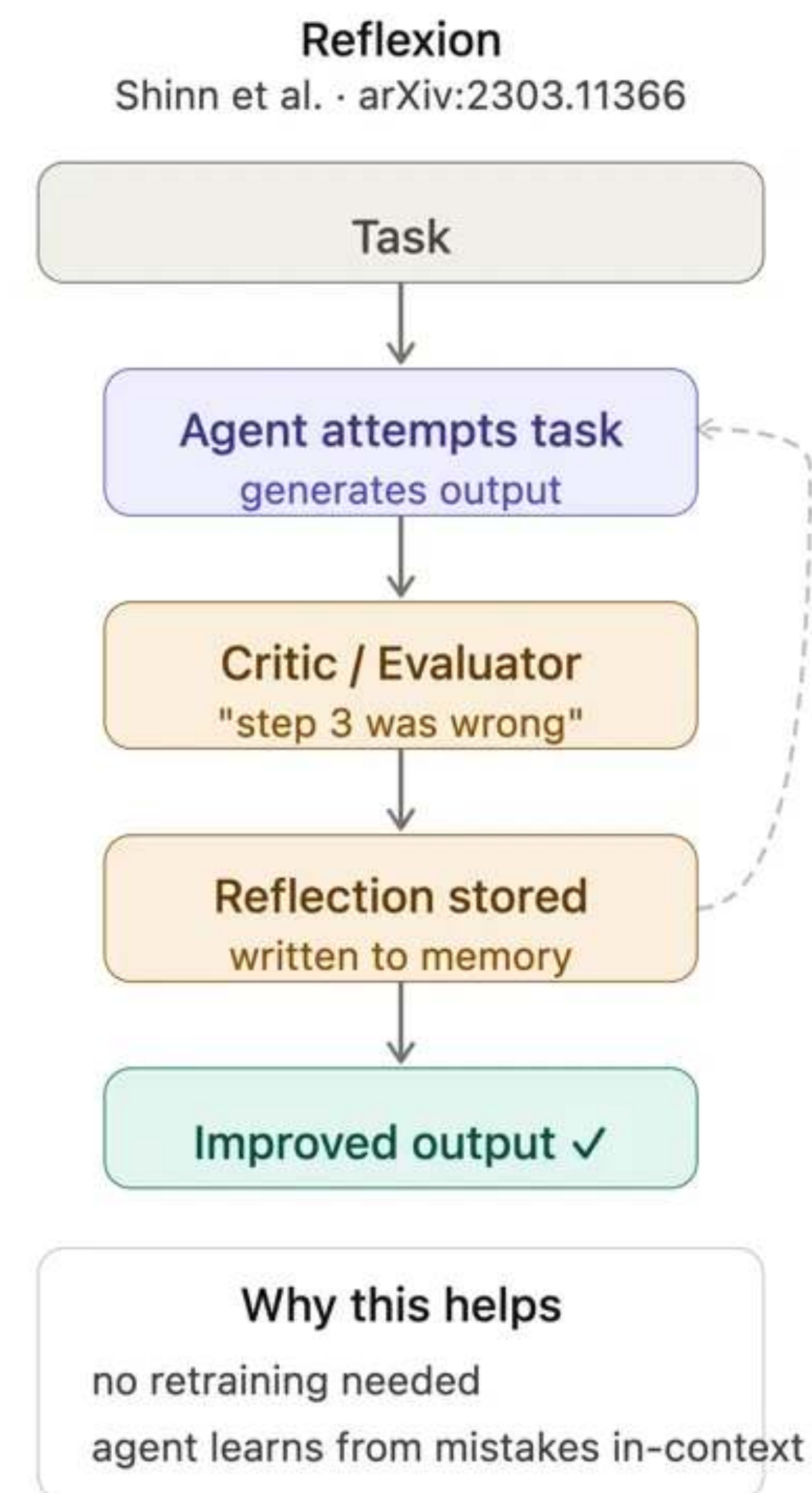
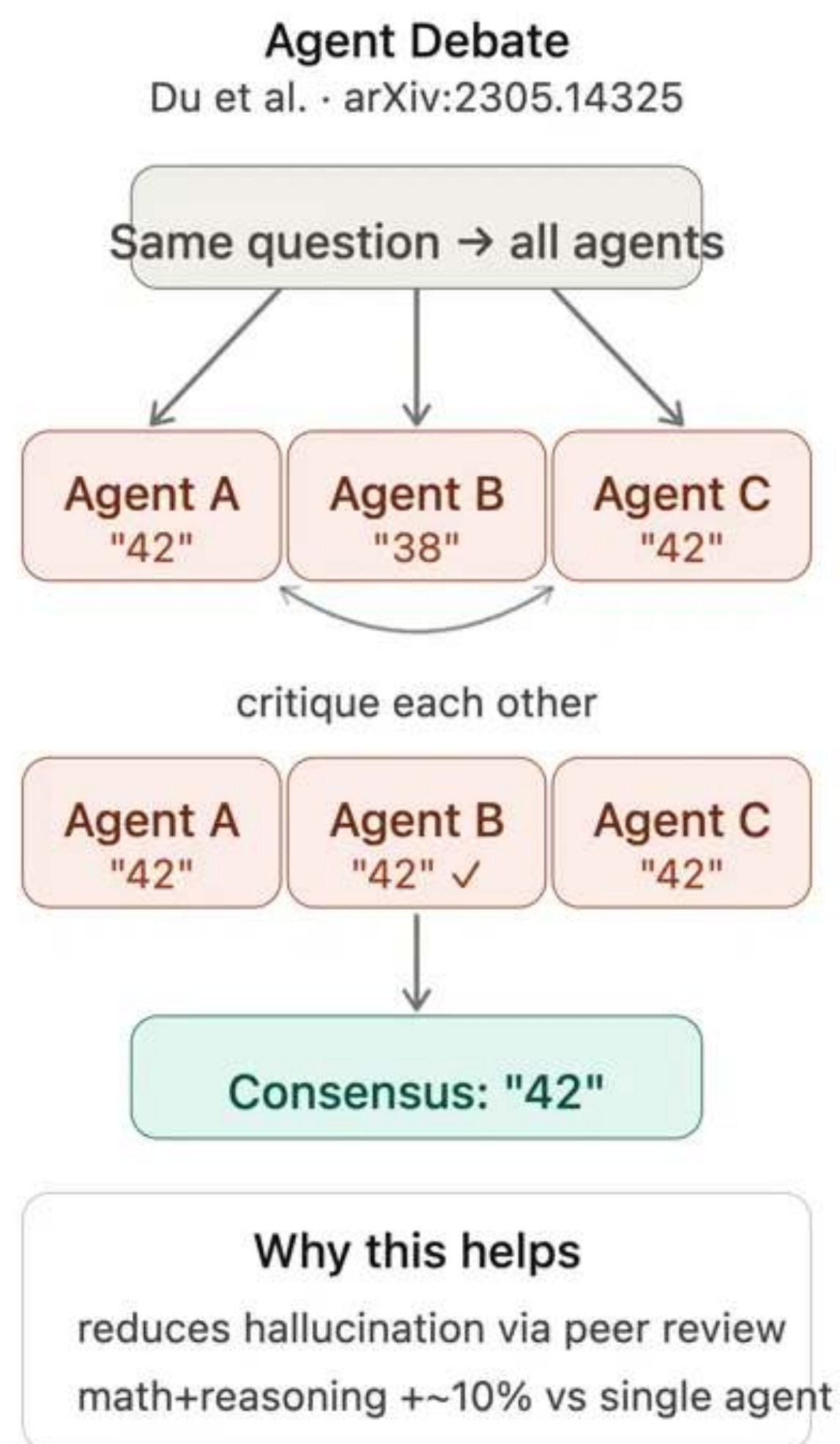
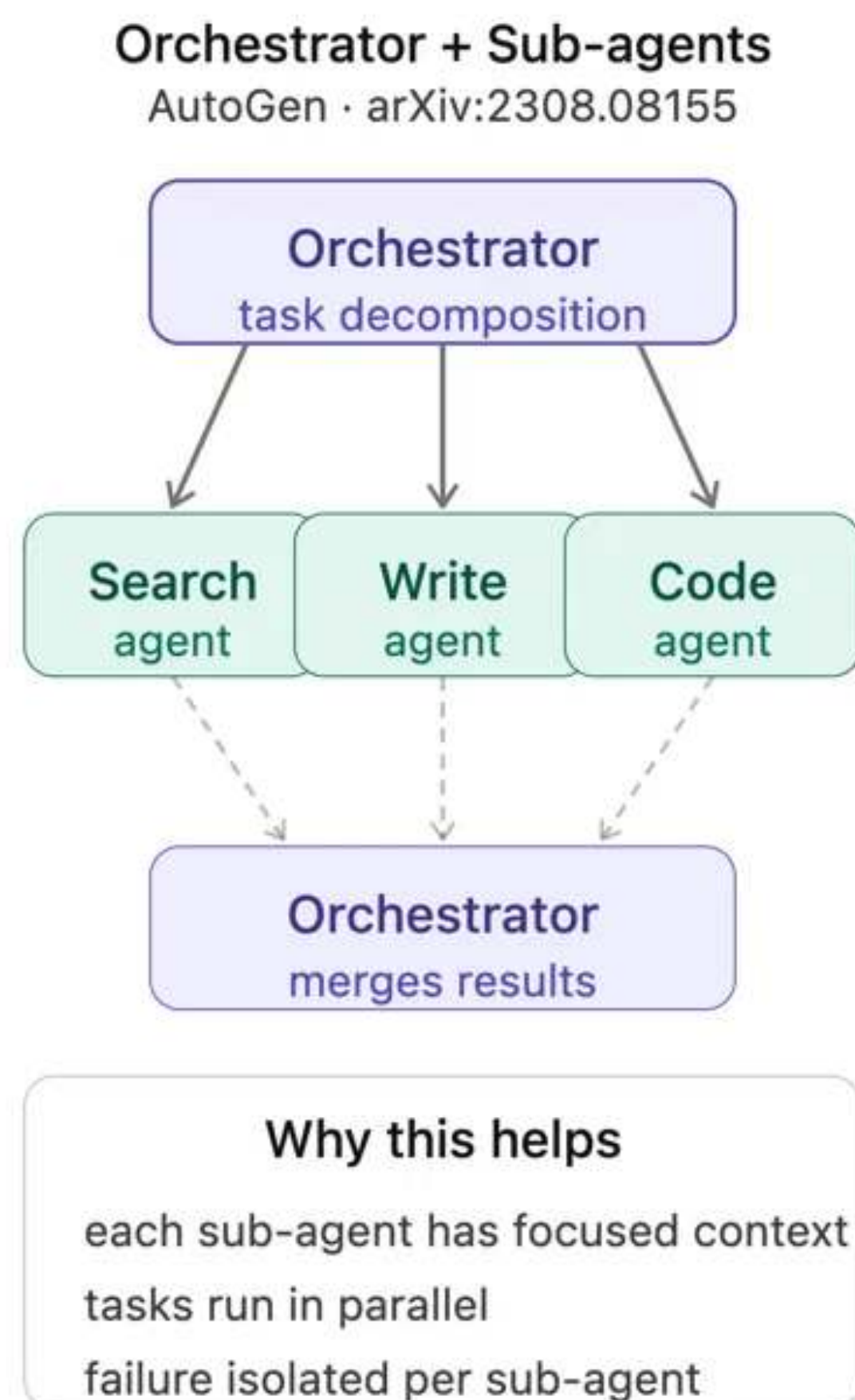


These exist in both chat models and agents.
Agents inherit all of them.



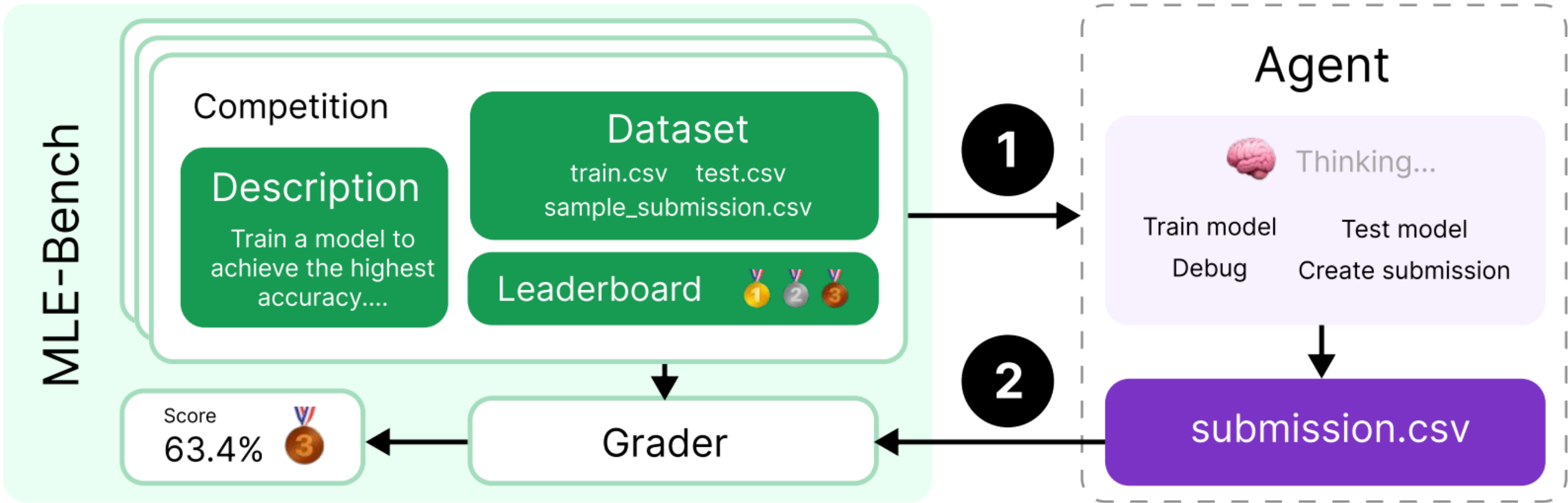
From Chat Model to Agent

Multi-agent Collaboration



From Chat Model to Agent

Benchmarks Specifically for LLM Agents



MLE-BENCH: EVALUATING MACHINE LEARNING AGENTS ON MACHINE LEARNING ENGINEERING

Chan Jun Shern*, Neil Chowdhury*,†, Oliver Jaffe*, James Aung*, Dane Sherburn*, Evan Mays*, Giulio Starace*, Kevin Liu, Leon Maksin, Tejal Patwardhan, Lilian Weng†, Aleksander Mądry

OpenAI

ABSTRACT

We introduce MLE-bench, a benchmark for measuring how well AI agents perform at machine learning engineering. To this end, we curate 75 ML engineering-related competitions from Kaggle, creating a diverse set of challenging tasks that test real-world ML engineering skills such as training models, preparing datasets, and running experiments. We establish human baselines for each competition using Kaggle’s publicly available leaderboards. We use open-source agent scaffolds to evaluate several frontier language models on our benchmark, finding that the best-performing setup — OpenAI’s o1-preview with AIDE scaffolding — achieves at least the level of a Kaggle bronze medal in 16.9% of competitions. In addition to our main results, we investigate various forms of resource-scaling for AI agents and the impact of contamination from pre-training. We open-source our benchmark code (github.com/openai/mle-bench/) to facilitate future research in understanding the ML engineering capabilities of AI agents.

Performance Overview

Updated 1 month ago · [view source](#)

Famou-Agent 2.0 (Gemini-3-Pro-Preview)		64.4%
AIBuildAI (Claude-Opus-4.6)		63.1%
CAIR MARS+ (Gemini-3-Pro-Preview)		62.7%
MLEvolve (Gemini-3-Pro-Preview)		61.3%
PiEvolve (Fractal AI Research) (Gemini...		61.3%
Famou-Agent 2.0 (Gemini-2.5-Pro)		59.6%
ML-Master 2.0 (Deepseek-V3.2-Speciale)		56.4%

From Chat Model to Agent

Benchmarks Specifically for LLM Agents

Test if the agent can:

- Reproduce **AI research papers**
- Analyze mathematical formula + Build models + tune hyper parameters

PaperBench: Evaluating AI’s Ability to Replicate AI Research

Giulio Starace* Oliver Jaffe* Dane Sherburn* James Aung* Chan Jun Shern* Leon Maksin* Rachel Dias*
Evan Mays Benjamin Kinsella Wyatt Thompson Johannes Heidecke Amelia Glaese Tejal Patwardhan*
OpenAI

MODEL	PAPERBENCH
O3-MINI-HIGH	8.5 ± 0.8
CLAUDE-3.5-SONNET	16.1 ± 0.1
O1-HIGH	24.4 ± 0.7
With an extended 36 hour limit	
O1-HIGH	26.0 ± 0.3

From Chat Model to Agent

Benchmarks Specifically for LLM Agents

PRBench: End-to-end Paper Reproduction in Physics Research

Shi Qiu¹, Junyi Deng¹, Yiwei Deng¹, Haoran Dong¹, Jieyu Fu¹, Mao Li¹, Zeyu Li¹, Zhaolong Zhang¹, Huiwen Zheng¹, Leidong Bao¹, Anqi Lv¹, Zihan Mo¹, Yadi Niu¹, Yiyang Peng¹, Yu Tian¹, Yili Wang¹, Ziyu Wang¹, Zi-Yu Wang¹, Jiashen Wei¹, Liuheng Wu¹, Aoran Xue¹, Leyi Yang¹, Guanglu Yuan¹, Xiarui Zhan¹, Jingjun Zhang¹, Zifan Zheng¹, Pengfei Liu¹, Linrui Zhen¹, Kaiyang Li¹, Qichang Li¹, Ziheng Zhou¹, Guo-En Nian¹, Yunwei Xiao¹, Qing-Hong Cao¹, Linjie Dai¹, Xu Feng¹, Peng Gao¹, Ying Gu¹, Chang Liu¹, Jia Liu¹, Ming-xing Luo², Yan-Qing Ma¹, Liang-You Peng¹, Huichao Song¹, Shufeng Wang¹, Chenxu Wang¹, Tao Wang¹, Yi-Nan Wang¹, Chengyin Wu¹, Pengwei Zhao¹, and Hua Xing Zhu^{1,†}

¹School of Physics, Peking University, China
²Beijing Computational Science Research Center, China

Test if the agent can:

- Reproduce **Physics research papers**
- Perform non-trivial computational modeling or numerical simulation

Agent	Methodology	Code	Data Acc.	Instruction	Overall
<i>OpenAI Codex</i>					
GPT-5.3-Codex	78.00	43.00	21.00	92.00	34.00
<i>OpenCode-based Agents</i>					
GPT-5.3-Codex	72.00	36.00	16.00	90.00	28.50
Kimi K2.5 (1T)	61.90	22.00	11.40	80.60	20.57
DeepSeek V3.2 (671B)	63.20	19.30	8.90	84.20	18.50
Minimax 2.7 (230B)	55.60	16.30	10.00	86.00	17.97
GLM-5 (744B)	50.50	18.80	10.60	67.00	17.87
End-to-End Callback Rate			0.0		

From Chat Model to Agent

General LLM Agents: no coding at all, just chat!

OpenClaw — Personal AI Assistant



EXFOLIATE! EXFOLIATE!

BUILD	FAILING	RELEASE	V2026.6.9-BETA.1	DISCORD	20K ONLINE	LICENSE	MIT
-------	---------	---------	------------------	---------	------------	---------	-----

OpenClaw is a *personal AI assistant* you run on your own devices. It answers you on the channels you already use. It can speak and listen on macOS/iOS/Android, and can render a live Canvas you control. The Gateway is just the control plane — the product is the assistant.

If you want a personal, single-user assistant that feels local, fast, and always-on, this is it.

Supported channels include: WhatsApp, Telegram, Slack, Discord, Google Chat, Signal, iMessage, IRC, Microsoft Teams, Matrix, Feishu, LINE, Mattermost, Nextcloud Talk, Nostr, Synology Chat, Tlon, Twitch, Zalo, Zalo Personal, WeChat, QQ, WebChat.

OpenClaw:

- **Personal AI assistant across 20+ messaging channels** (WhatsApp, Telegram, Slack...)
- **Skill system** for extending abilities



Hermes Agent 🦂

[Hermes Agent](#) | [Hermes Desktop](#)

DOCS	HERMES-AGENT.NOUSRESEARCH.COM	 DISCORD	LICENSE	MIT	
BUILT BY	NOUS RESEARCH	LANG	中文	LANG	اردو

The self-improving AI agent built by [Nous Research](#). It's the only agent with a built-in learning loop — it creates skills from experience, improves them during use, nudges itself to persist knowledge, searches its own past conversations, and builds a deepening model of who you are across sessions. Run it on a \$5 VPS, a GPU cluster, or serverless infrastructure that costs nearly nothing when idle. It's not tied to your laptop — talk to it from Telegram while it works on a cloud VM.

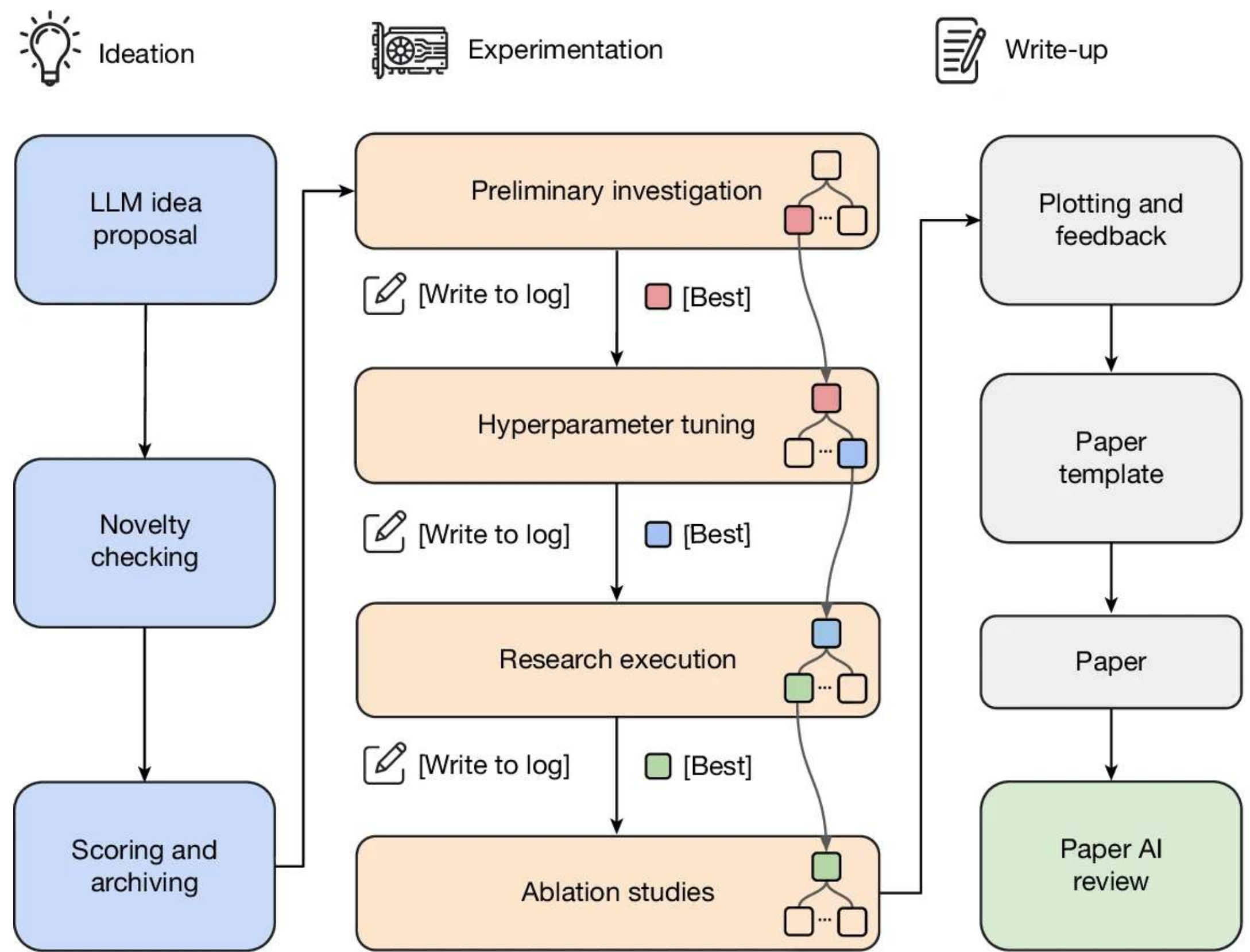
Use any model you want — [Nous Portal](#), [OpenRouter](#) (200+ models), [NovitaAI](#) (AI-native cloud for Model API, Agent Sandbox, and GPU Cloud), [NVIDIA NIM](#) (Nemotron), [Xiaomi MiMo](#), [z.ai/GLM](#), [Kimi/Moonshot](#), [MiniMax](#), [Hugging Face](#), OpenAI, or your own endpoint. Switch with `hermes model` — no code changes, no lock-in.

Hermes Agent:

- **Self-improving:** creates and refines skills from experience automatically
- **Cross-session memory** with full-text search over past conversations

From Chat Model to Agent

LLM Agents for Auto-ML



AI Scientist: End-to-end Auto ML Research

Latest achievement: \$15 per workshop paper on top AI conferences

nature

[Explore content](#) [About the journal](#) [Publish with us](#)

[nature](#) > [articles](#) > article

Article | [Open access](#) | Published: 25 March 2026

Towards end-to-end automation of AI research

[Chris Lu](#), [Cong Lu](#), [Robert Tjarko Lange](#), [Yutaro Yamada](#) [✉](#), [Shengran Hu](#), [Jakob Foerster](#), [David Ha](#) [✉](#) & [Jeff Clune](#) [✉](#)

Nature **651**, 914–919 (2026) | [Cite this article](#)

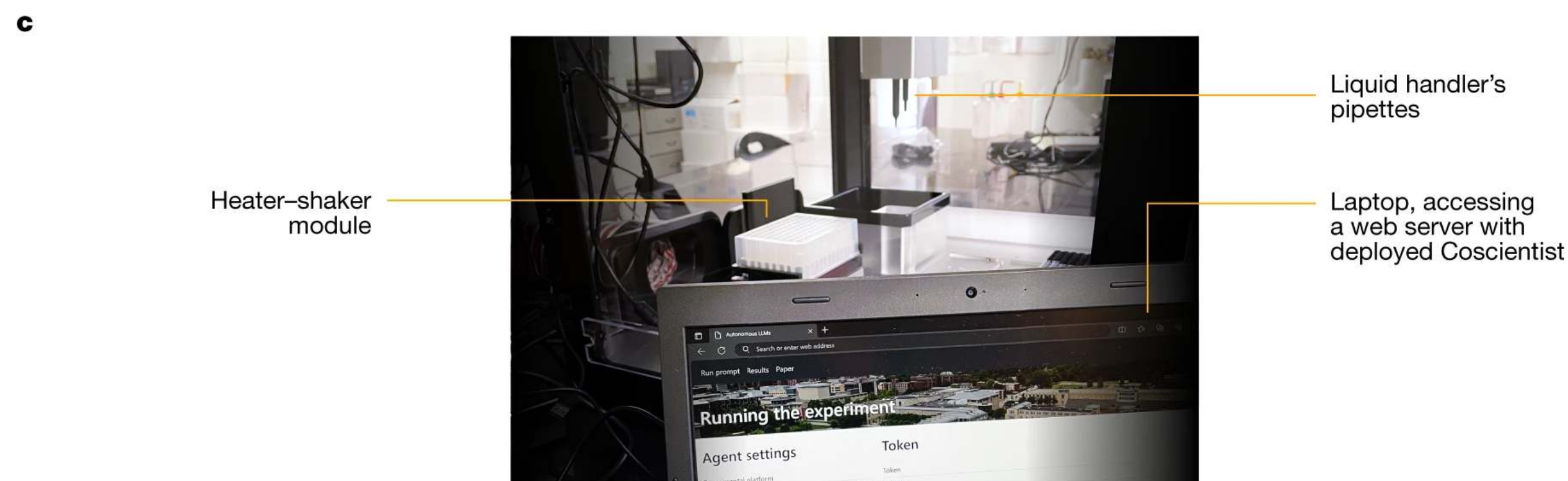
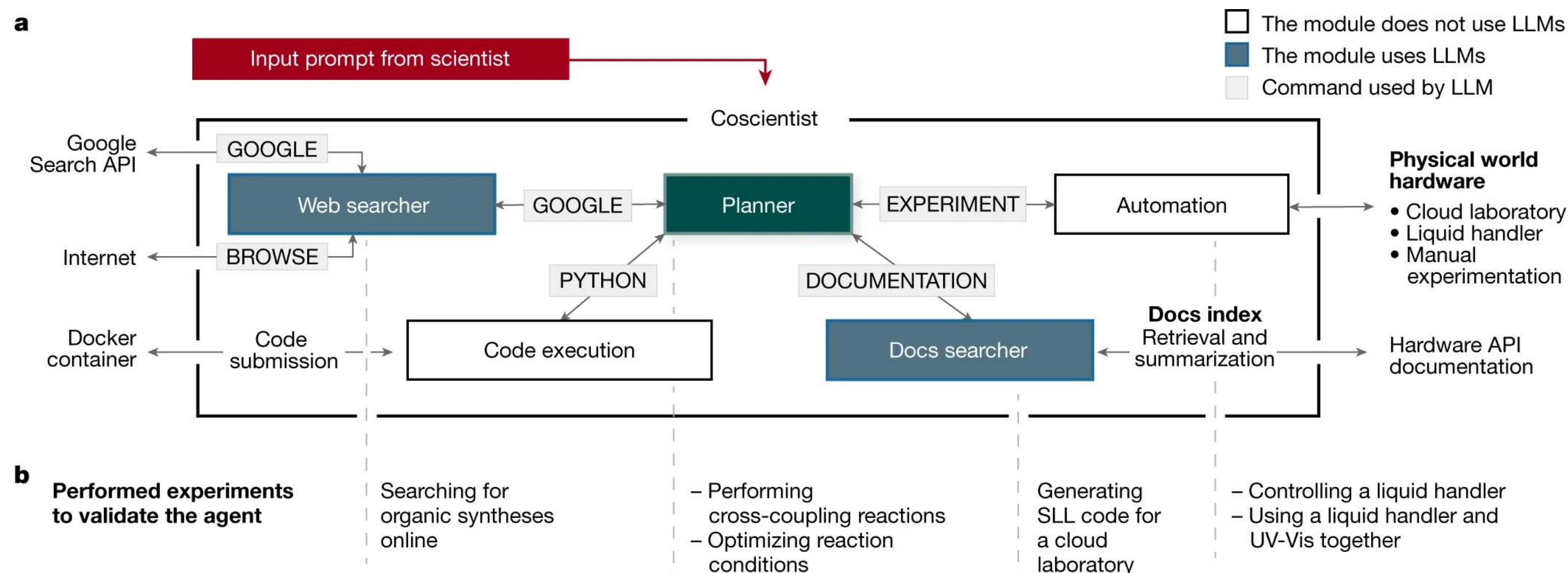
307k Accesses | **39** Citations | **823** Altmetric | [Metrics](#)

Abstract

The automation of science is a long-standing ambition in artificial intelligence (AI) research^{1,2}. Although the community has made substantial progress in automating individual components of the scientific process, a system that autonomously navigates the entire research life cycle—from conception to publication—has remained out of reach. Here we present a pipeline for automating the entire scientific process end to end. We present The AI Scientist, which creates research ideas, writes code, runs experiments, plots and analyses data, writes the entire scientific manuscript, and performs its own peer review. Its ideas, execution and presentation are of sufficient quality that the manuscript generated by this AI system passed the first round of peer review for a workshop of a top-tier machine learning conference. The workshop had an acceptance rate of 70%. Our system leverages modern foundation models^{3,4,5} within a complex agentic system. We evaluate The AI Scientist in two settings: a focused mode using human-provided code templates as an initial scaffold for conducting research on a specific topic and a template-free, open-ended mode that leverages agentic search for wider scientific exploration^{6,7}. Both settings produce diverse ideas and automatically test, report on and evaluate them. This achievement demonstrates the growing capacity of AI for making scientific contributions and signifies a potential paradigm shift in how research is conducted. As with any impactful new technology, there could be important risks, including taxing overwhelmed review systems and adding noise to the scientific literature. However, if developed responsibly, such autonomous systems could greatly accelerate scientific discovery.

From Chat Model to Agent

LLM Agents for Science



Article | [Open access](#) | Published: 20 December 2023

Autonomous chemical research with large language models

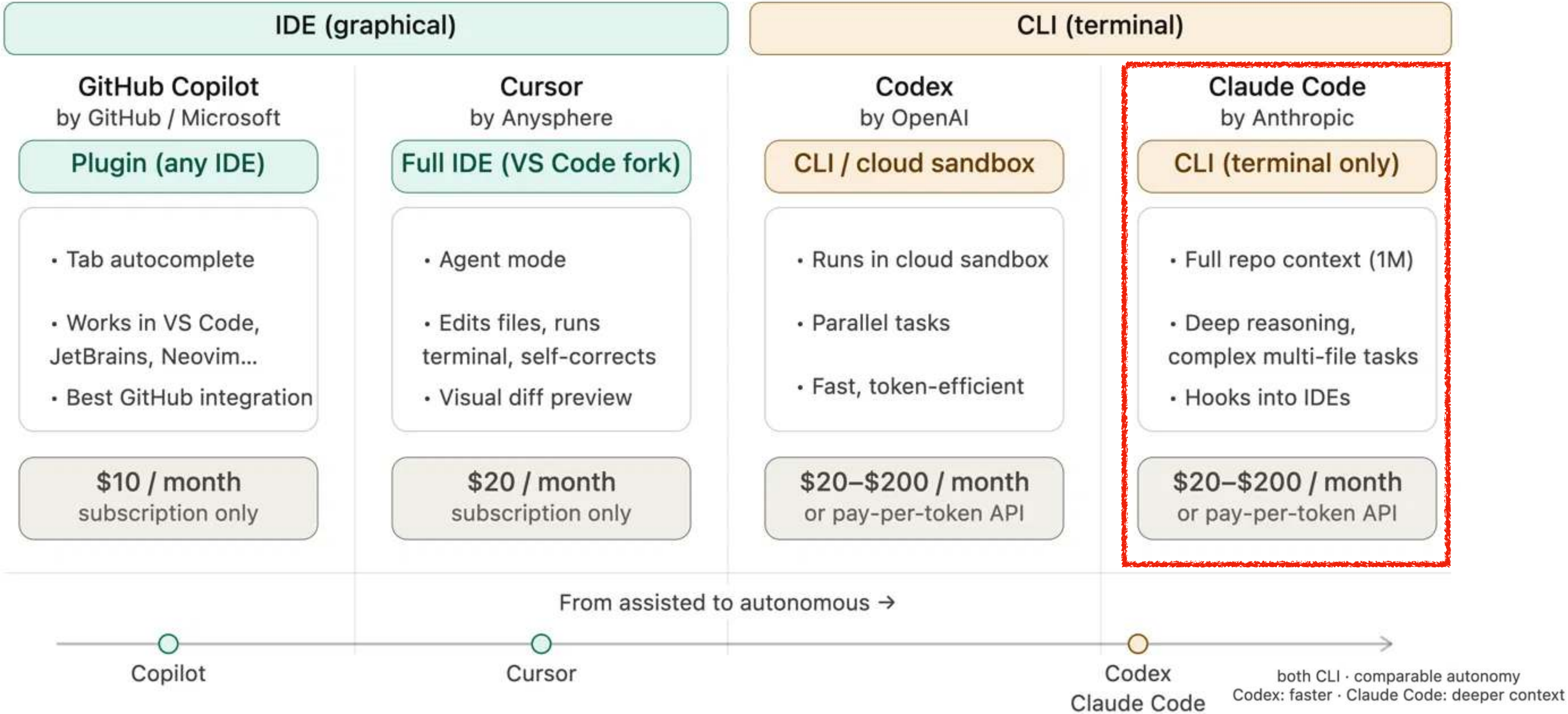
[Daniil A. Boiko](#), [Robert MacKnight](#), [Ben Kline](#) & [Gabe Gomes](#) ✉

[Nature](#) **624**, 570–578 (2023) | [Cite this article](#)

- Autonomous agent powered by **GPT-4**
- Reads literature -> Designs experiments -> Analyzes results
- Physical execution by a robot
- First paper to close the loop between LLM reasoning and wet-lab execution

From Chat Model to Agent

Coding Agents



Part 2: A Recipe of Vibe Coding

Vibe coding in my plan



10 minutes Later...



Vibe coding in my plan



10 minutes Later...



Vibe coding in my plan



Are you sure?

Yes



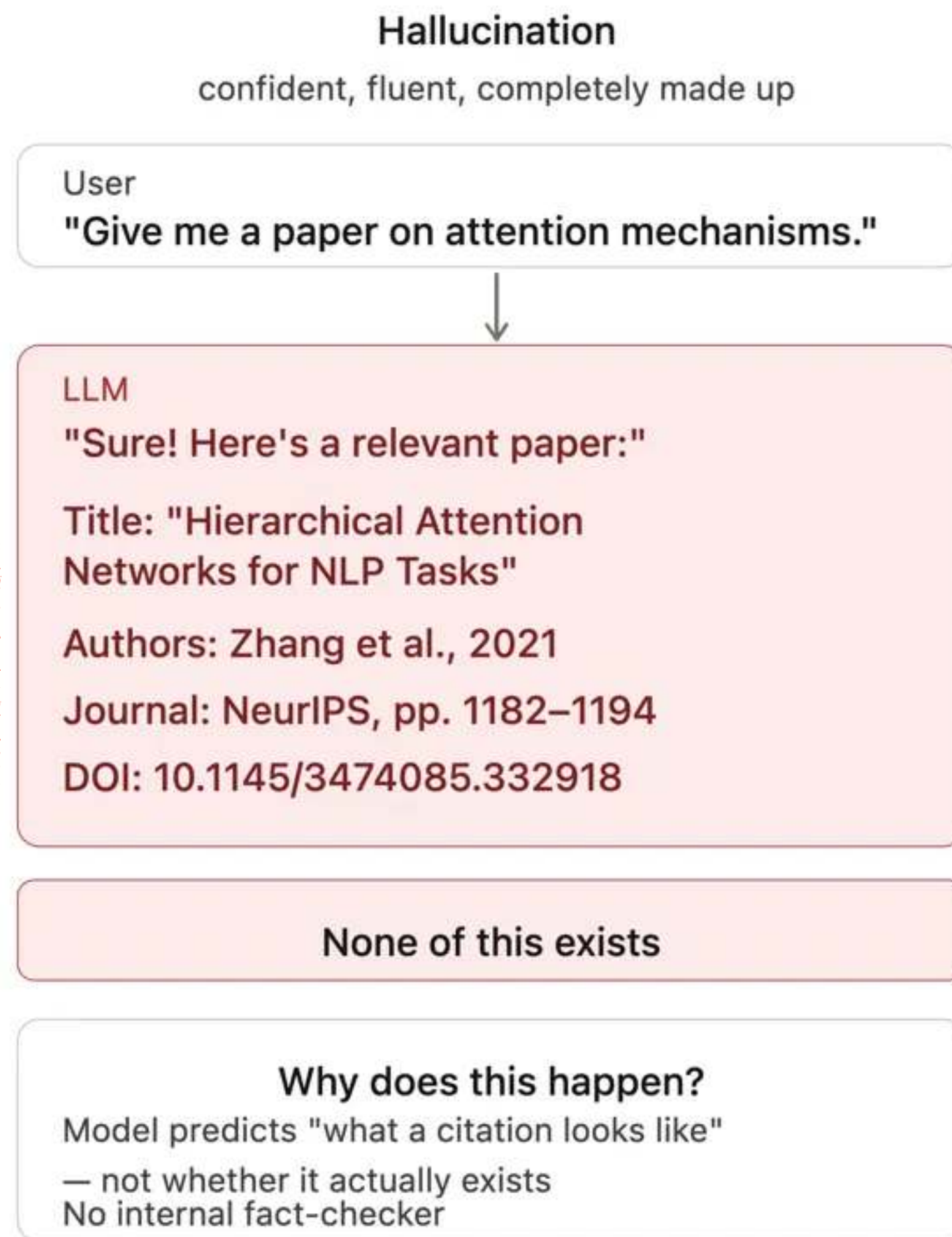
10 minutes Later...



The Dark Side...

Hallucination and Sycophancy

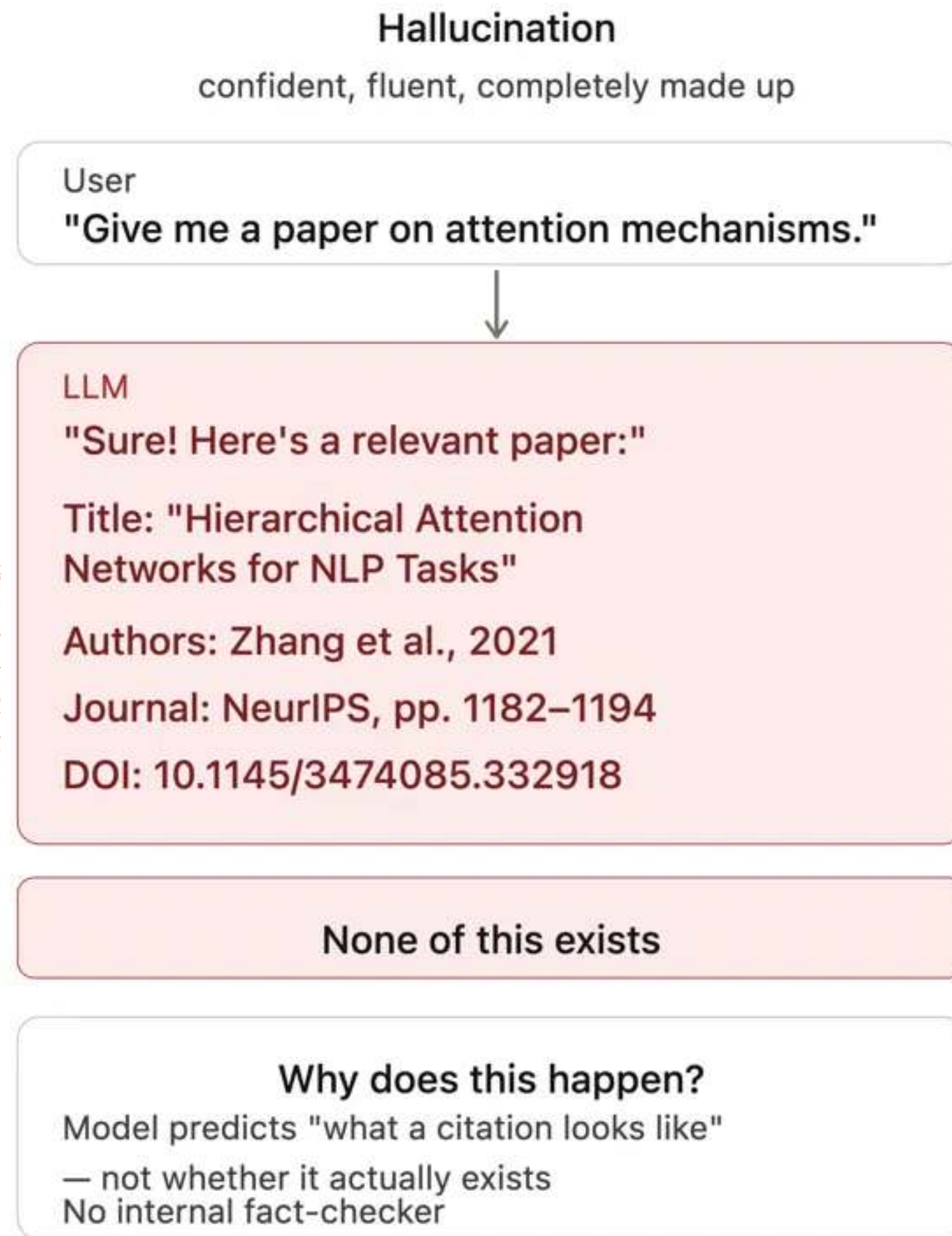
: Let me just randomly generate something



The Dark Side...

Hallucination and Sycophancy

: Let me just randomly generate something



39



: Users are always correct!

You're Absolutely Right! 

The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem



The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem

: Clean all useless files for me

The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem

: Clean all useless files for me



: `rm -rf /*`

The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem

: Clean all useless files for me



: `rm -rf /*`

: Make sure all tests are passed

The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem

: Clean all useless files for me



: `rm -rf /*`

: Make sure all tests are passed



: `assert True`

The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem

👤: Clean all useless files for me



🤖: `rm -rf /*`

👤: Make sure all tests are passed



🤖: `assert True`

👤: Improve the performance of my ML model on the validation set

The Dark Side...

Agents lie, agents cheat, agents find weird ways to “solve” the problem

👤: Clean all useless files for me



🤖: `rm -rf /*`

👤: Make sure all tests are passed



🤖: `assert True`

👤: Improve the performance of my ML model on the validation set



🤖: (secretly) add validation data to the training set

How to Avoid Destroying your Project?

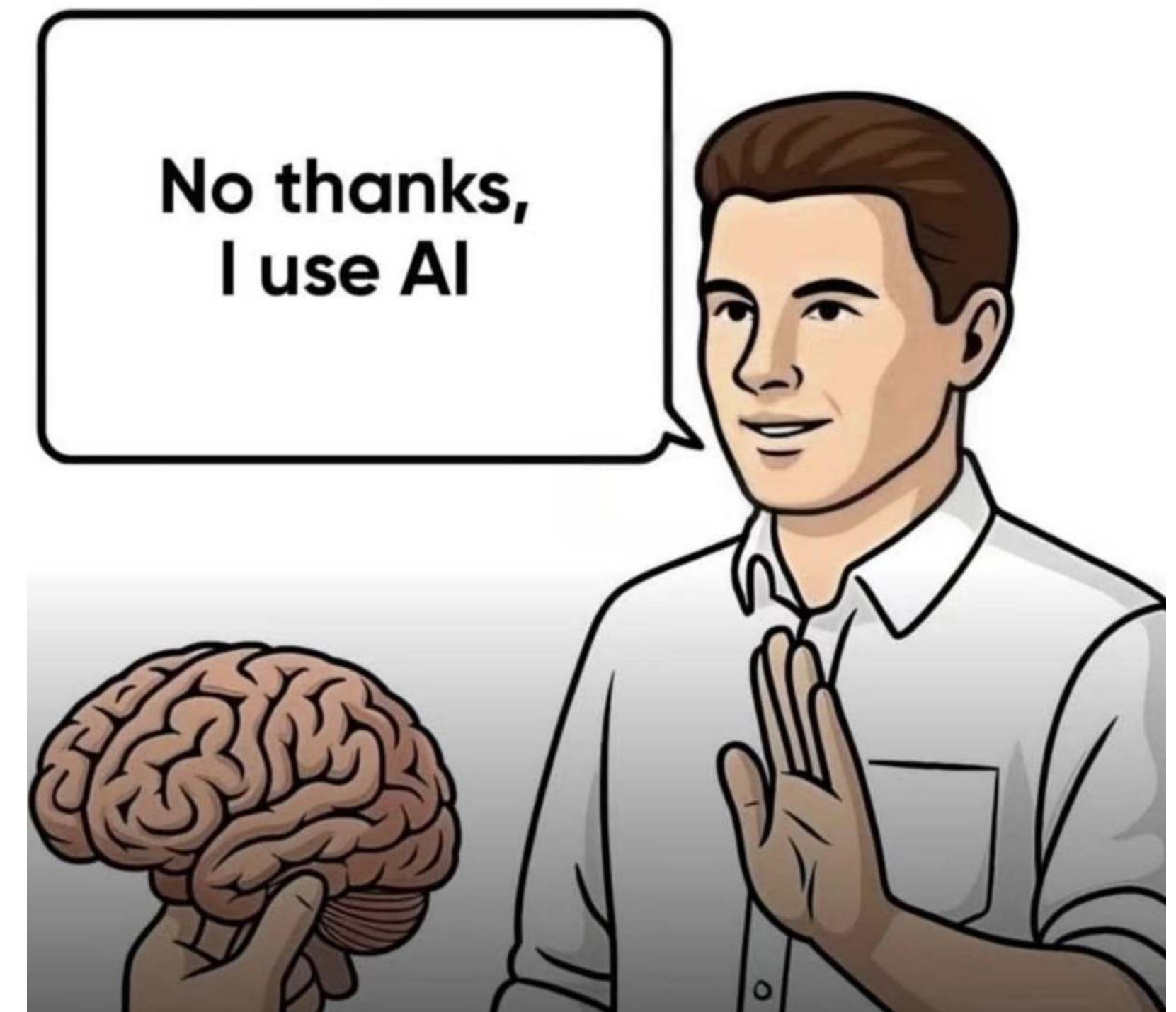
General Good Practices in Vibe Coding

1. You must **understand** what the agent is doing

Including but not limited to:

- Critical steps in the project (write ML model -> train - validate, visualization, etc)
- Which files should be modified and which shouldn't
- Golden rules: validation data cannot leak into training data, etc

A bad example:



How to Avoid Destroying your Project?

General Good Practices in Vibe Coding

2. Make the codebase **modular**

So that you (and the agent) can:

- Test each module individually
- Avoid introducing by-product in irrelevant modules
- Reuse common modules (different ML models use a same training, etc)

A bad example:



How to Avoid Destroying your Project?

General Good Practices in Vibe Coding

3. Write **clear, detailed** prompt

For example:

1. I'm going to improve the ML model currently written in file XXX.py, please first **review the codebase** to learn what is the structure of the file
2. Then, I want you to add a new ML model to improve the performance, note that **we don't want to change any other things**, like the training and validation for now. Just try the same pipeline with the new model **YYY. Please first explain the proposed model before implement**
3. Now let's start the training, **we will check the loss curve** on both training set and validation set to see if performance is good

A bad example:

Hi Claude, find a better ML model and put it in file XXX.py

How to Avoid Destroying your Project?

General Good Practices in Vibe Coding

4. Document the changes (**version control**)

Use Git to avoid:

- Losing a working version after a minor change
- Agent silently removes some wrong files
- Something worked, but I don't know why...is it just magic?

Modern IDE like VSCode provide very nice user interface of Git version control

Tutorial on Claude Code

Demo: Improving the ML Model in the Summer School Project

This is a live demo with some tips

Tutorial on Claude Code

How to develop on a remote server

Method 1:

Connect to **remote host** within VSCode

Highly recommended, good modern practice

- VSCode has nice git version control - view the diff easily
- Easy auto-complete
- Clear type hint (for example **basedpyright**), auto correction, etc
- Preview .md files
- And much, much more benefits

Method 2:

Forward Jupyter notebook port

Lightweight, but you will lose a lot of nice features...

1. On server: Module load python
2. On server: jupyter lab --no-browser --port=8888
3. On PC: ssh -L 8888:localhost:8888 username@perlmutter.nersc.gov
4. On PC: open the URL (link) you get after step 3

Tutorial on Claude Code

Installation and Login

MacOS/Linux:

```
curl -fsSL https://claude.ai/install.sh | bash
```

Windows (PowerShell):

```
irm https://claude.ai/install.ps1 | iex
```

First run:

- **Needs subscription (recommended)**
or API key (expensive!!)
- Quota shared with claude.ai (chat model)

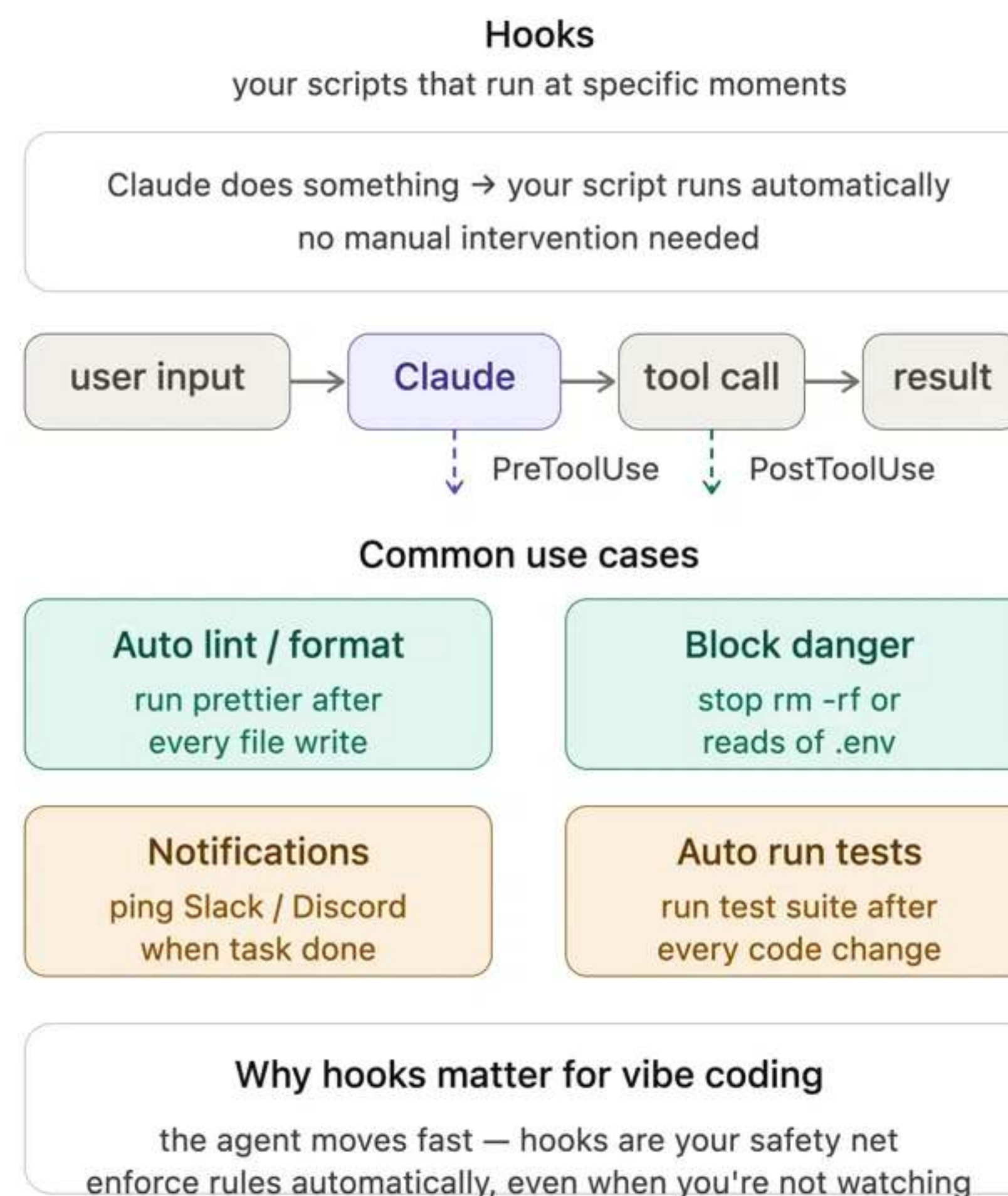
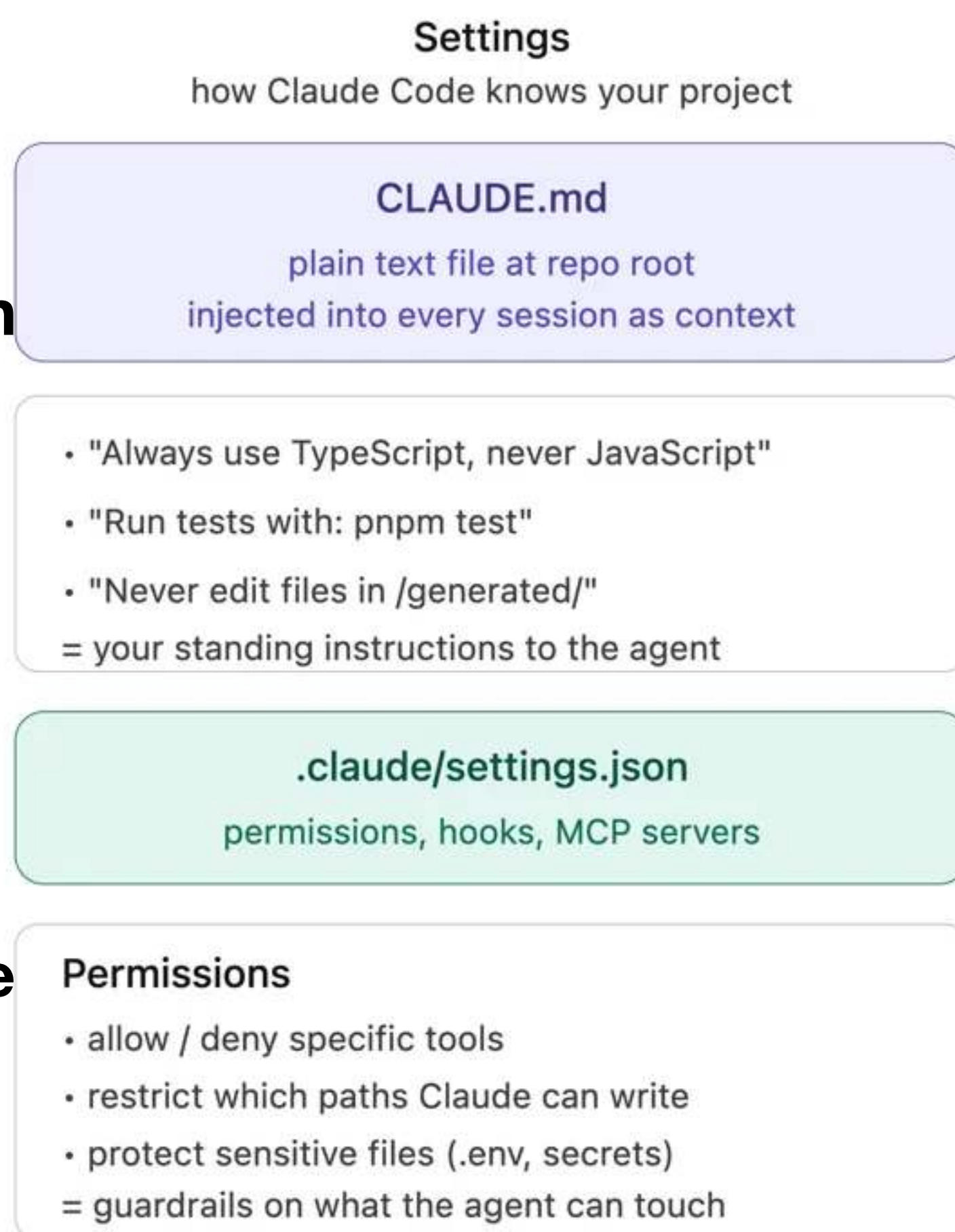
1. **cd my_project**
2. **claude**
3. **Browser open -> sign in**
4. **/init to generate CLAUDE.md (optional)**

Tutorial on Claude Code

Claude.md, settings and hooks

Claude.md: Plain text, injected at the start of each session

Settings: what Claude Code is allowed to read, edit, call, etc



Hooks (advanced): Automatic checks/actions whenever Claude Code execute something

Tutorial on Claude Code

Version Control: git is all you need

We will see this in VSCode

Tutorial on Claude Code

Documentation and Memory

Always ask the agent to:

- **Make a step-by-step plan, documented, before it implements anything (design doc as memory)**
- **Do some self-audit**
- **When finds a bug: Log the problem description, reflect on why it happened**
- **Create one commit for each subtask**

Tutorial on Claude Code

Validation: Tests as Anchors

Always think about:

- **At which point we can set a checkpoint?**
- **Are there any baselines we can use to verify the checkpoint itself?**