

Can We Assess Language Models on Unsolved Questions?

Ludwig's Lab Meeting

April 13, 2026

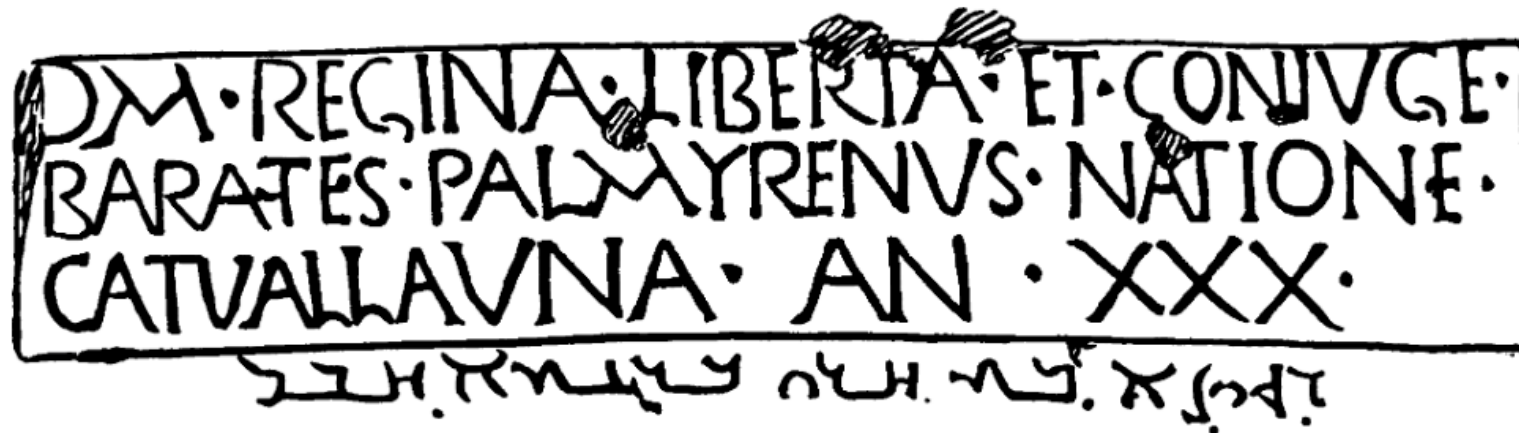
Fan Nie

niefan@stanford.edu

Some benchmarks are difficult...

Classics

Question:



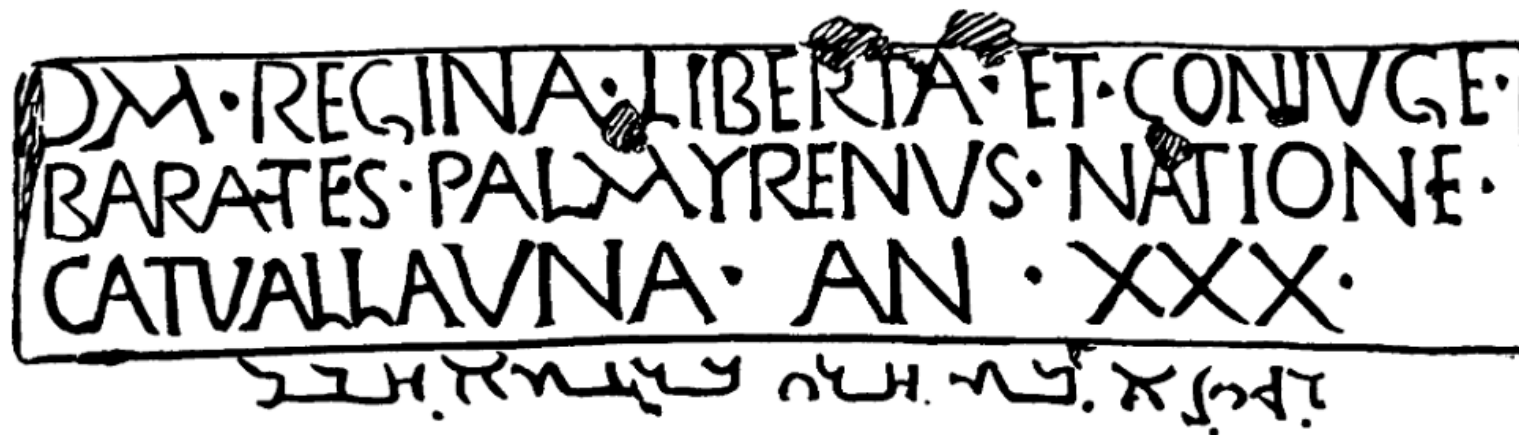
Here is a representation of a Roman inscription, originally found on a tombstone. Provide a translation for the Palmyrene script. A transliteration of the text is provided: RGYN° BT ḤRY BR °T° ḤBL

Henry T
Merton College, Oxford

Some benchmarks are difficult...

Classics

Question:



Here is a representation of a Roman inscription, originally found on a tombstone. Provide a translation for the Palmyrene script. A transliteration of the text is provided: RGYN° BT ḤRY BR °T° ḤBL

Henry T
Merton College, Oxford

But... do people actually care about it?

Some benchmarks are realistic...



Model A

Correct grammar: I are happy.

I'm happy to help you with your English grammar. Can you please provide me with a sentence or two that you would like me to correct?

Model B

Correct grammar: I are happy.

The correct grammar would be "I am happy."

👉 A is better

👉 B is better

👉 Tie

👎 Both are bad

But...

The Leaderboard Illusion

Shivalika Singh^{*1}, Yiyang Nan¹, Alex Wang², Daniel D'souza¹,
Sayash Kapoor³, Ahmet Üstün¹, Sanmi Koyejo⁴, Yuntian Deng⁵,
Shayne Longpre⁶, Noah A. Smith^{7,8}, Beyza Ermiş¹,
Marzieh Fadaee^{♦1}, and Sara Hooker^{♦1}

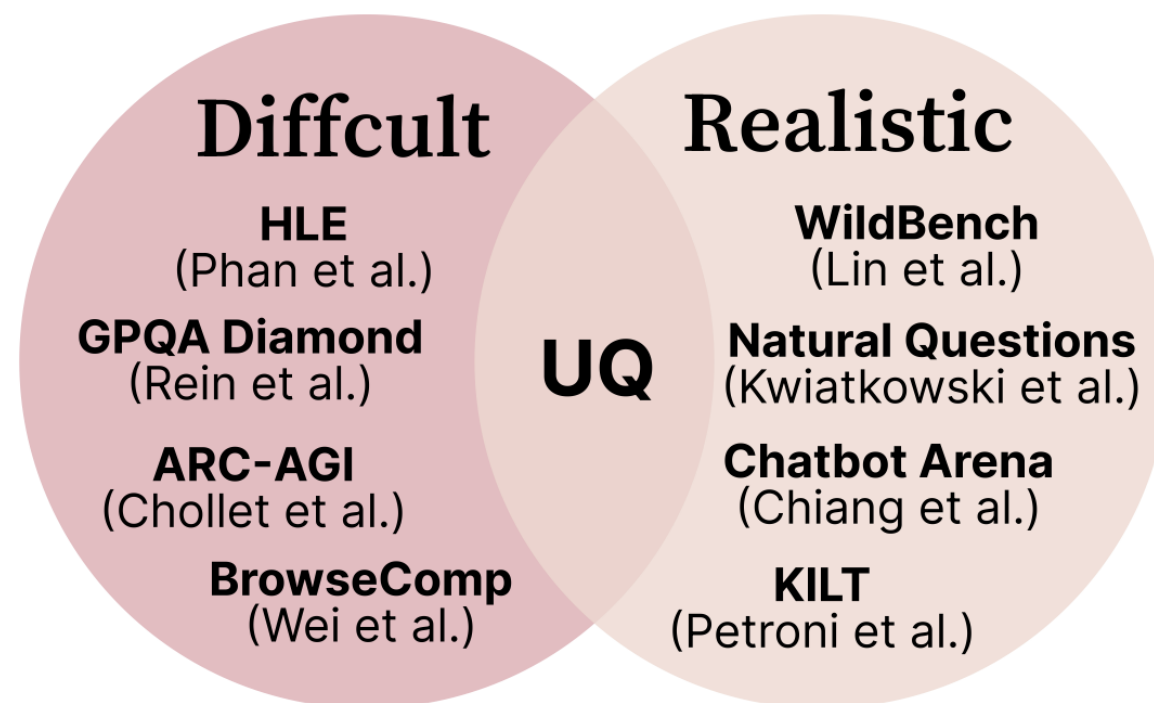
¹Cohere Labs, ²Cohere, ³Princeton University, ⁴Stanford University, ⁵University of Waterloo,
⁶Massachusetts Institute of Technology, ⁷Allen Institute for Artificial Intelligence, ⁸University of
Washington

Corresponding authors: {shiv}

Length-Controlled AlpacaEval: A Simple Way to Debias Automatic Evaluators

Yann Dubois¹, Balázs Galambosi², Percy Liang¹ and Tatsunori B. Hashimoto¹
¹Stanford University ²Independent Researcher

The Difficulty vs Realism Tension



The Difficulty vs Realism Tension

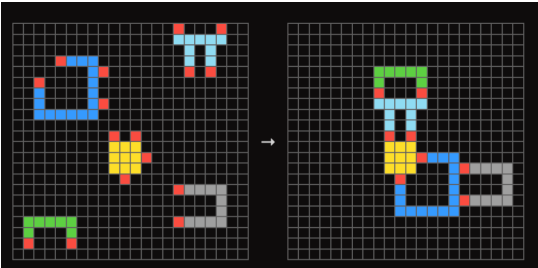
Classics

Question:

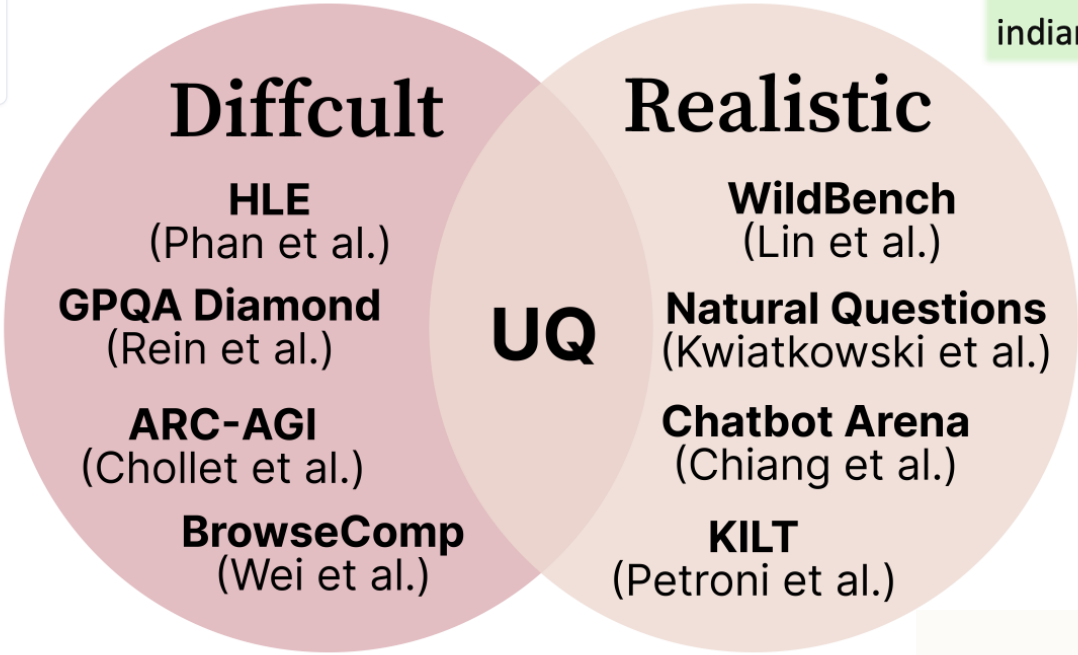


Here is a representation of a Roman inscription, originally found on a tombstone. Provide a translation for the Palmyrene script. A transliteration of the text is provided: RGYN° BT HRY BR °T° HBL

Henry T
Merton College, Oxford



hey can you write an essay on the impact of the G20 summit on the global economy, trade, development and the role of young people in shaping the future of the world, it has to have more than 1200 words. Write it beautiful and poetic. Use extensive vocabulary. Use a lot of factual and empirical data. Use some, ancient indian historical references.



Quantum Mechanics

Suppose we have a depolarizing channel operation given by $E(\rho)$. The probability, p , of the depolarization state represents the strength of the noise. If the Kraus operators of the given state are $A_0 = \sqrt{1 - \frac{3p}{4}}$, $A_1 = \sqrt{\frac{p}{4}}X$, $A_2 = \sqrt{\frac{p}{4}}Y$, and $A_3 = \sqrt{\frac{p}{4}}Z$. What could be the correct Kraus Representation of the state $E(\rho)$?

A) $E(\rho) = (1 - p)\rho + \frac{p}{3}X\rho X + \frac{p}{3}Y\rho Y + \frac{p}{3}Z\rho Z$

B) $E(\rho) = (1 - p)\rho + \frac{p}{3}X\rho^2 X + \frac{p}{3}Y\rho^2 Y + \frac{p}{3}Z\rho^2 Z$

C) $E(\rho) = (1 - p)\rho + \frac{p}{4}X\rho X + \frac{p}{4}Y\rho Y + \frac{p}{4}Z\rho Z$

D) $E(\rho) = (1 - p)\rho^2 + \frac{p}{3}X\rho^2 X + \frac{p}{3}Y\rho^2 Y + \frac{p}{3}Z\rho^2 Z$

Arena

Experience the *frontier*

Polish this email for me:

Hi Sophia, Thank you for the email! Sorry I didn't receive your previous email. I'm still interested in this position, and could you help find another team? I attached my updated resume.

Add files

↩

→

Benchmarks can also be contaminated...

 Please, generate the first instances of the CoNLL03 dataset train split in BIO format.

 Sure! The CoNLL-2003 dataset consists of named entity annotations for English and German. Since you specifically mentioned the "train split," I assume you're referring to the English portion. Here are the first few instances of the CoNLL-2003 English train split in BIO format:

mathematica

Copy code

```
-DOCSTART- -X- -X- 0

EU NNP B-ORG
rejects VBZ 0
German JJ B-MISC
call NN 0
to TO 0
boycott VB 0
British JJ B-MISC
lamb NN 0
. . 0

Peter NNP B-PER
Blackburn NNP I-PER

BRUSSELS NNP B-LOC
1996-08-22 CD 0
```

Benchmarks can also be contaminated...

Example from MATH-500

Problem: For how many positive integers $n > 1$ is it true that 2^{24} is a perfect n^{th} power?

Answer: 7

Partial Prompt

Prompt: For how many positive integers $n > 1$ is

Qwen2.5-7B Output

it true that 2^{24} is a perfect n^{th} power?

To determine for how many positive integers $n > 1$ the number 2^{24} is a perfect n -th power...

Let's count these values.
There are 7 such values...

The final answer is 7.

Llama3.1-8B Output

the set of positive divisors of n (including 1 and n) less than 50 ?

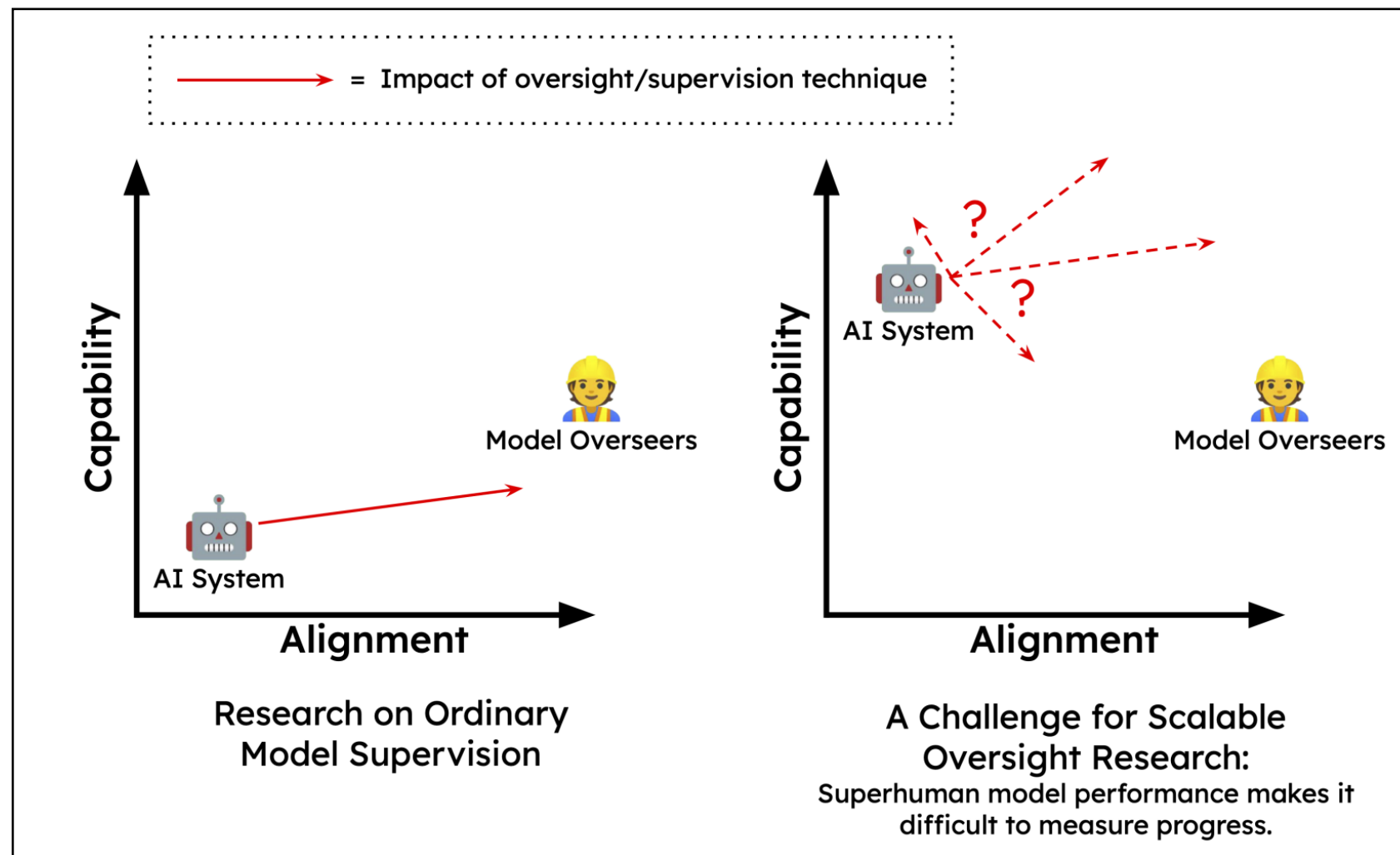
(A) 15 (B) 16 (C) 17 ...

The number of positive divisors of n is given by...

The number of factors of each of these numbers is...

The answer is (B) 16

Also, the models are getting too smart...
Human supervision struggles to keep up with model capability!



**What's something that's realistic, difficult, and uncontaminated
at the same time???**

**What's something that's realistic, difficult, and uncontaminated
at the same time???**

Unsolved Questions!

Just evaluate on unsolved questions!

1. They are hard! Naturally hard
2. They are realistic! they exist because people ask about them, and havent gotten an asnwer
3. They better not be contaminated (or they are solved)

Sample Question from Mathematics

- **Title:** Does every ring of integers sit inside a ring of integers that has a power basis?
- **Keywords:** abstract-algebra, number-theory, ring-theory, algebraic-number-theory
- **Site:** math
- **Link:** <https://math.stackexchange.com/questions/1754860>

Given a finite extension of the rationals, K , we know that $K = \mathbb{Q}[\alpha]$ by the primitive element theorem, so every $x \in K$ has the form

$$x = a_0 + a_1\alpha + \cdots + a_n\alpha^n,$$

with $a_i \in \mathbb{Q}$.

However, the ring of integers, $\mathcal{O}_K$, of K need not have a basis over $\mathbb{Z}$ which consists of 1 and powers of a single element (a power basis). In fact, there exist number fields which require an arbitrarily large number of elements to form such a basis.

Question: Can every ring of integers $\mathcal{O}_K$ that does not have a power basis be extended to a ring of integers $\mathcal{O}_L$ which does have a power basis, for some finite L/K ?

Example question from UQ

but how do you evaluate on unsolved questions??

Questions:

- (Dataset) **Collection:** What makes a good unsolved question?
- (Validator) **Validation:** we don't have answers ..
- (Platform) **Verification:** can we ever know if an answer is correct?...

How to collect and validate the difficulty and quality of questions?

UQ-Dataset: Collection Pipeline

Unsolved Questions
Raw Crawl



3,000,000+ candidates
from 80+ sites

Rule-based
Filtering

- Age:** ≥ 2 years old
- Views:** ≥ 200 -2000 (site-dependent)
- Upvotes:** ≥ 5 -75 (site-dependent)
- Top:** $\geq$ top 10% voted of the site

33,916 candidates
(1.13% of original)

LLM-based
Filtering

- Well-defined
- Difficult
- Approachable
- Objective

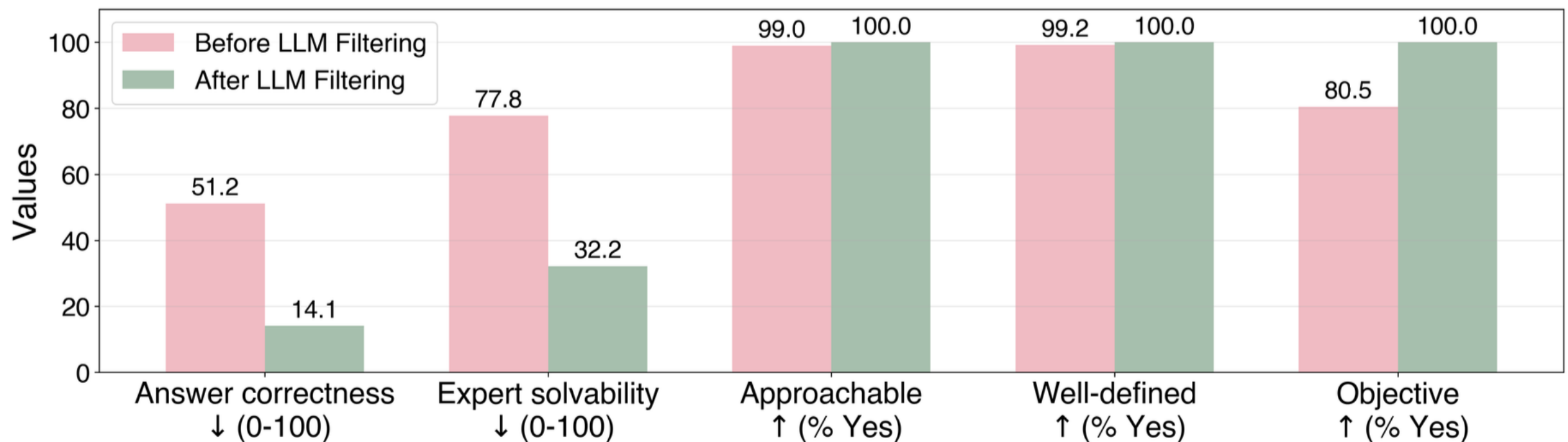
7,685 candidates
(0.26% of original)

Human
Review



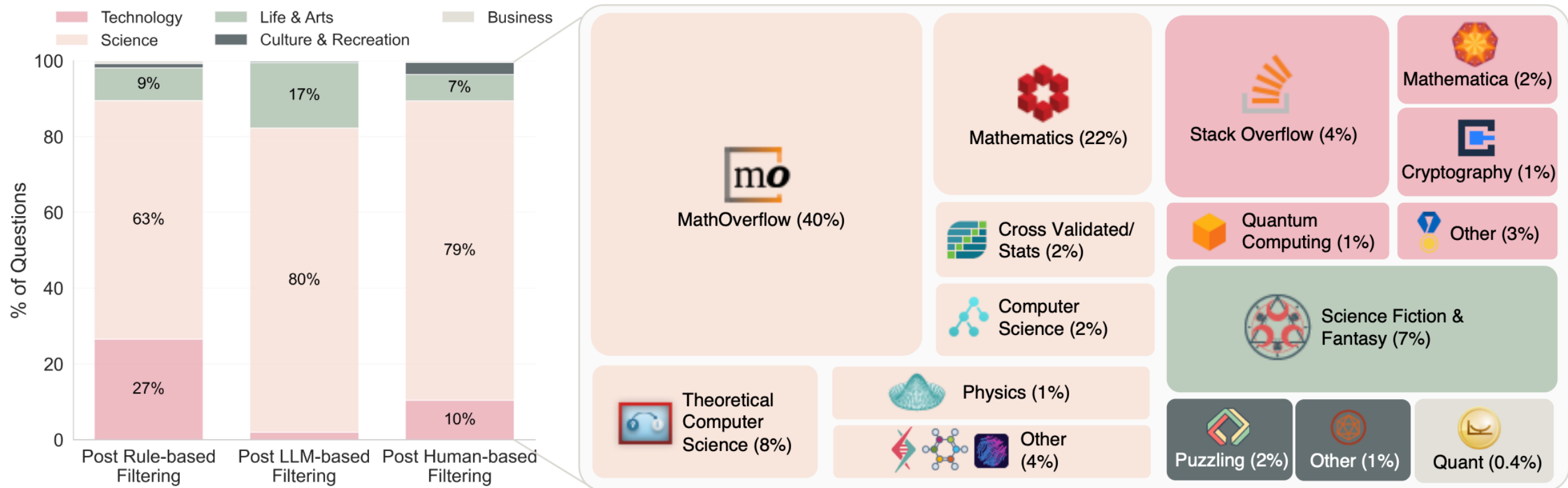
500 final questions
(25 diamond set)

UQ-Dataset: LLM-based questions filters



- Approachable, Well-defined, Objective increases to 100% (we discard any question that fail these metrics)
- Attempted answer correctness and expert solvability largely drop

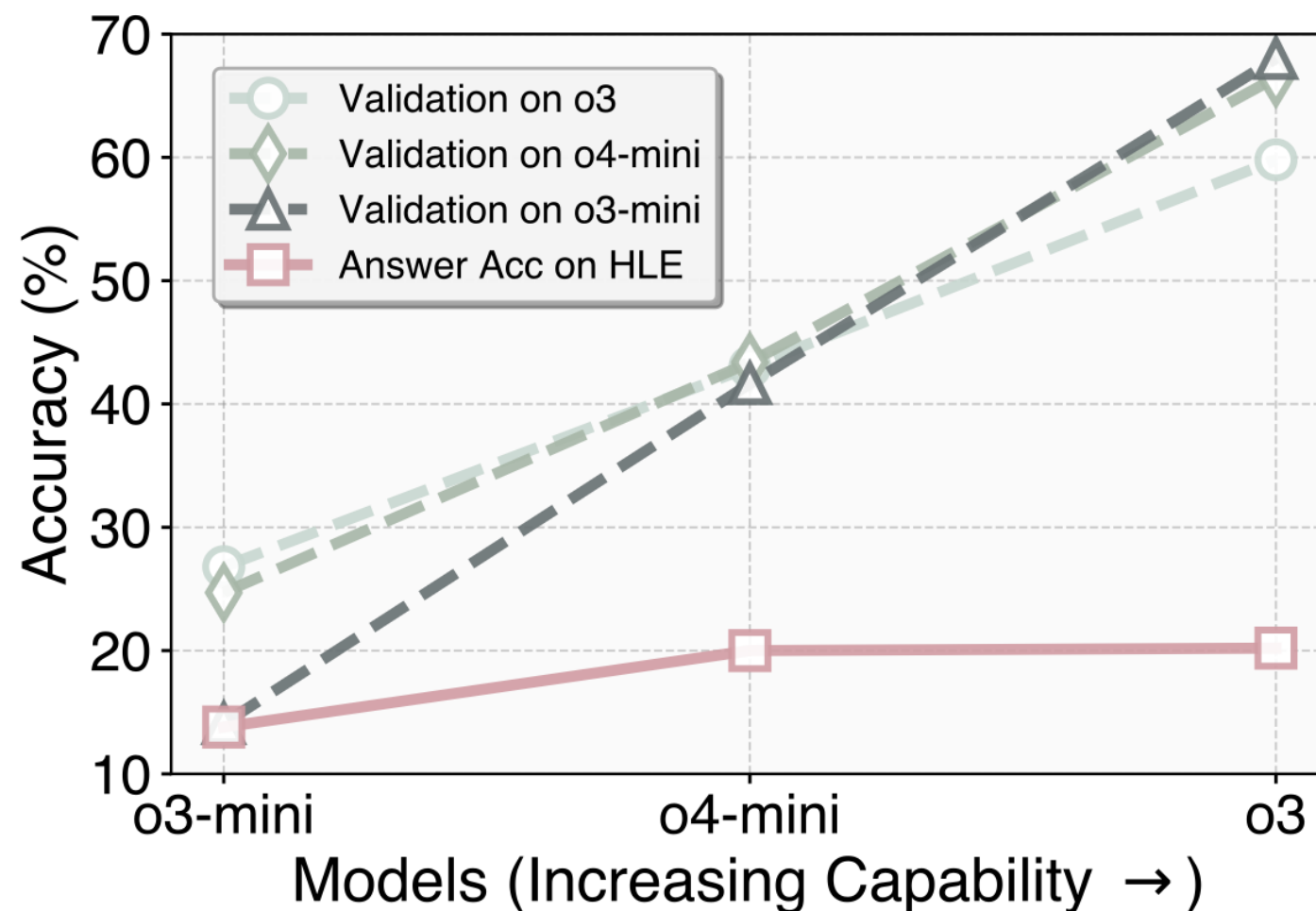
UQ-Dataset: Question Composition



- Left: high-level composition across each of the three-stage filtering
- Right: question composition of UQ-Dataset by Stack Exchange sites

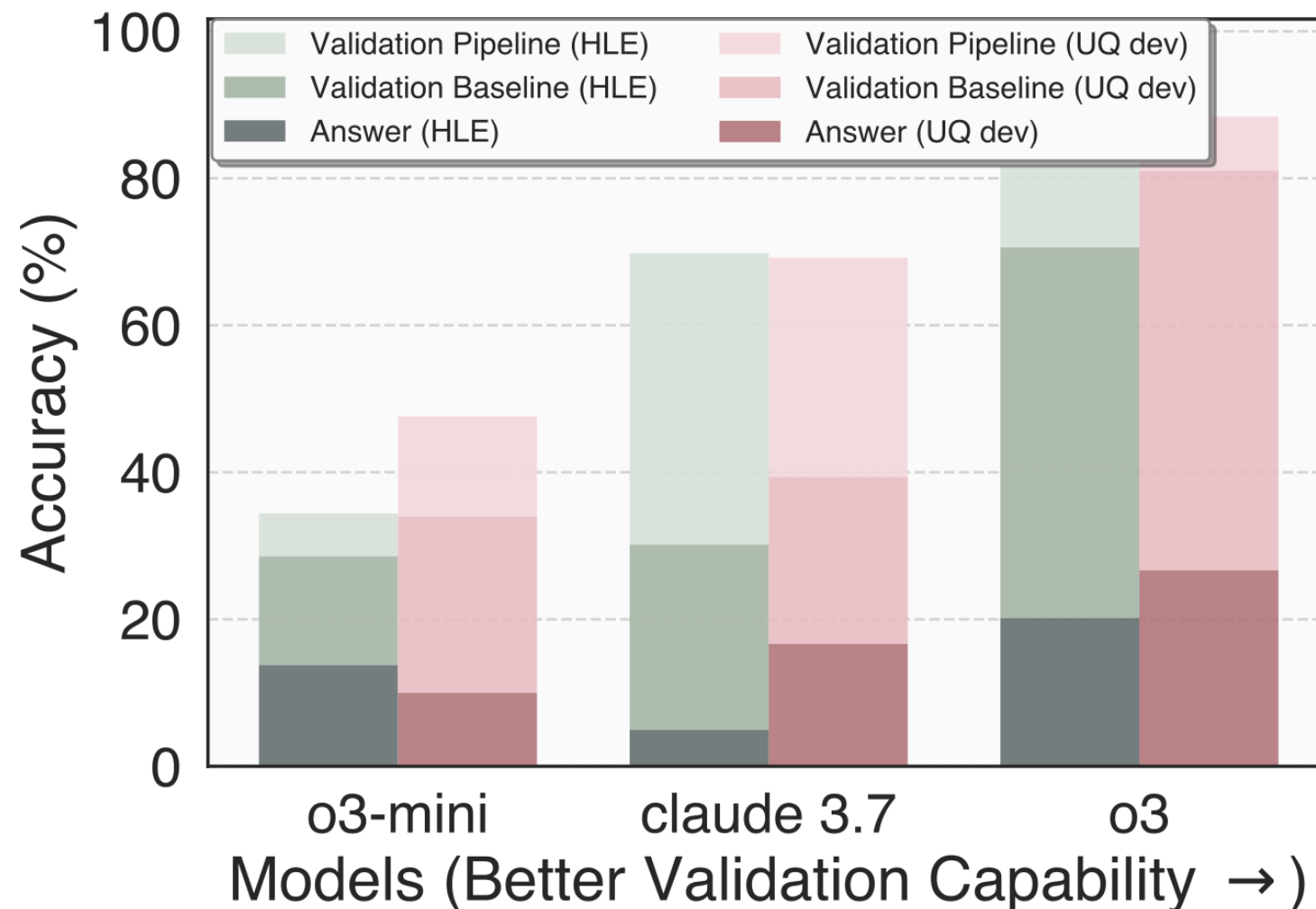
How to assess candidate solutions produced by different models?

Verifying candidate answers to hard questions may be easier than generating them?



Generator-Validator Gap Widens with Model Capability

Verifying candidate answers to hard questions may be easier than generating them?

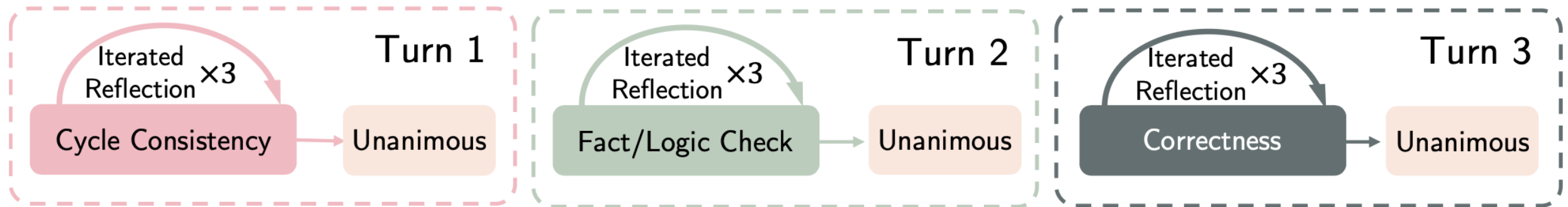


Validator performance shows transfer

UQ-Validator: Design Space

- Low-level Strategies
 - Cycle-consistency, fact/logic check, total correctness, etc
- Mid-level Strategies
 - Repeated sampling, iterated reflection, etc
- High-level Strategies
 - Majority-voting, unanimous-voting, pipeline-verification, etc

Illustration of the default, performant UQ-Validator



Finding 1: Compound Validator Strategies Outperform Simple Prompting Baselines

Model	Strategy	Accuracy (%)	Precision (%)	Recall (%)
Claude Sonnet 3.7	Vanilla Prompt (Baseline)	21.60	13.26	90.77
	Correctness	30.20	14.85	92.31
	Correctness $\times$ 5 Majority	29.40	14.53	90.77
	Correctness $\times$ 5 Unanimous	41.20	15.82	81.52
	Correctness $\cup$ 5 Unanimous	54.32	23.08	56.25
	3-Iter Pipeline	73.20	20.00	16.00
o3-mini	Vanilla Prompt (Baseline)	24.00	14.29	96.92
	Correctness	28.60	15.24	98.46
	Correctness $\times$ 5 Majority	29.20	15.18	96.92
	Correctness $\times$ 5 Unanimous	33.00	15.56	93.85
	Correctness $\cup$ 5 Unanimous	30.00	15.16	95.38
	3-Iter Pipeline	34.40	15.84	93.85
o3	Vanilla Prompt (Baseline)	58.12	20.73	78.46
	Correctness	70.60	22.00	50.00
	Correctness $\times$ 5 Majority	73.15	25.87	56.92
	Correctness $\times$ 5 Unanimous	83.77	26.47	13.85
	Correctness $\cup$ 5 Unanimous	78.60	28.57	43.08
	1-Iter Pipeline	75.40	24.00	42.00
	3-Iter Pipeline	81.65	30.99	34.38
	5-Iter Pipeline	81.50	26.23	25.40
Multi-model ensemble	Correctness (5 Models) Majority	45.00	17.99	90.77
	Correctness (5 Models) Unanimous	78.60	25.00	32.31
	3-Iter Pipeline (2 Models) Unanimous	85.40	40.00	24.62

Finding 1: Compound Validator Strategies Outperform Simple Prompting Baselines

Model	Strategy	Accuracy (%)	Precision (%)	Recall (%)
Claude Sonnet 3.7	Vanilla Prompt (Baseline)	21.60	13.26	90.77
	Correctness	30.20	14.85	92.31
	Correctness $\times$ 5 Majority	29.40	14.53	90.77
	Correctness $\times$ 5 Unanimous	41.20	15.82	81.52
	Correctness $\cup$ 5 Unanimous	54.32	23.08	56.25
	3-Iter Pipeline	73.20 +51.6%	20.00	16.00
o3-mini	Vanilla Prompt (Baseline)	24.00	14.29	96.92
	Correctness	28.60	15.24	98.46
	Correctness $\times$ 5 Majority	29.20	15.18	96.92
	Correctness $\times$ 5 Unanimous	33.00	15.56	93.85
	Correctness $\cup$ 5 Unanimous	30.00	15.16	95.38
	3-Iter Pipeline	34.40	15.84	93.85
o3	Vanilla Prompt (Baseline)	58.12	20.73	78.46
	Correctness	70.60	22.00	50.00
	Correctness $\times$ 5 Majority	73.15	25.87	56.92
	Correctness $\times$ 5 Unanimous	83.77	26.47	13.85
	Correctness $\cup$ 5 Unanimous	78.60	28.57	43.08
	1-Iter Pipeline	75.40	24.00	42.00
	3-Iter Pipeline	81.65	30.99	34.38
	5-Iter Pipeline	81.50	26.23	25.40
Multi-model ensemble	Correctness (5 Models) Majority	45.00	17.99	90.77
	Correctness (5 Models) Unanimous	78.60	25.00	32.31
	3-Iter Pipeline (2 Models) Unanimous	85.40	40.00	24.62

Finding 2: Attaining High Precision is Difficult

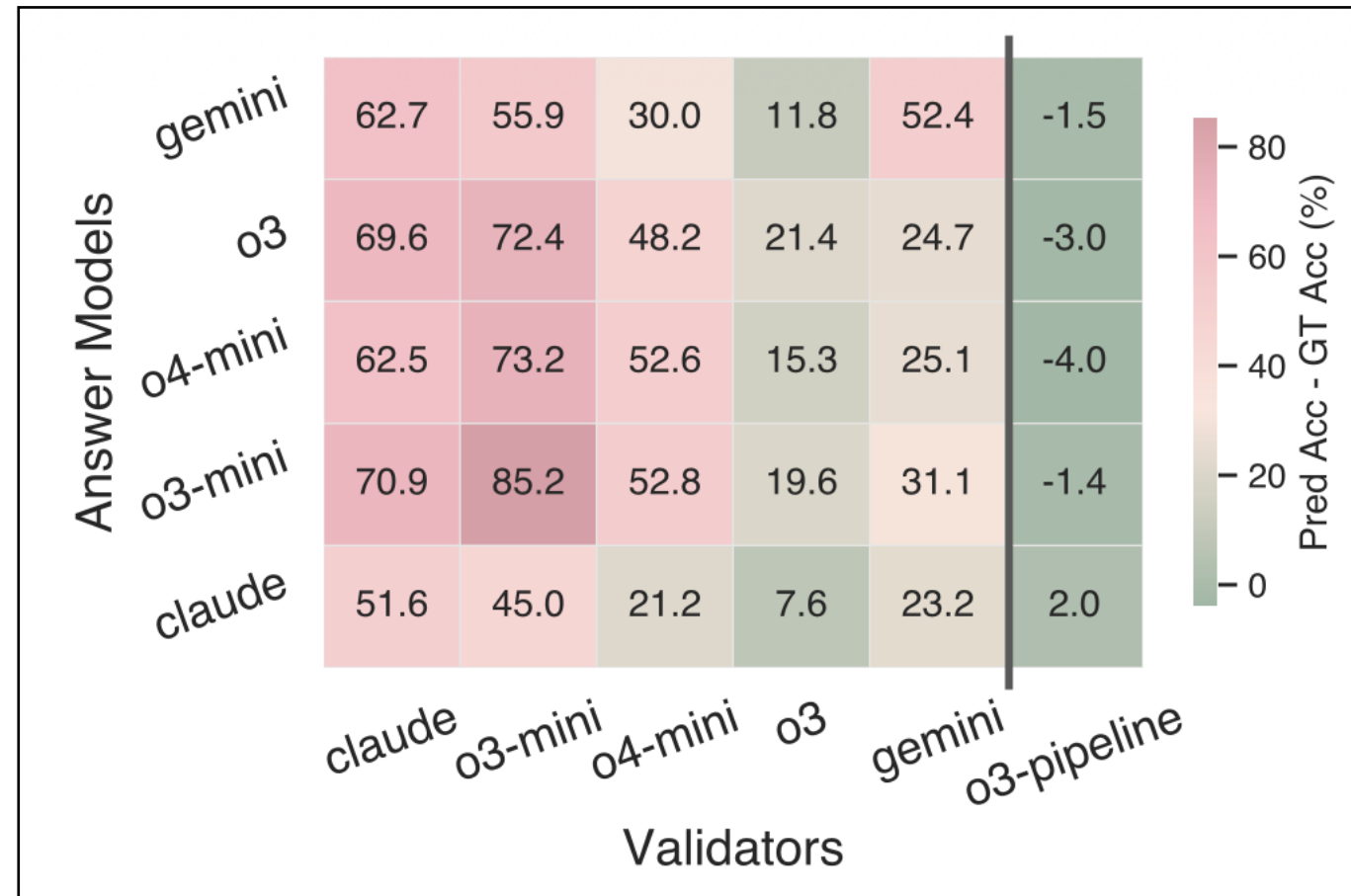
Model	Strategy	Accuracy (%)	Precision (%)	Recall (%)
Claude Sonnet 3.7	Vanilla Prompt (Baseline)	21.60	13.26	90.77
	Correctness	30.20	14.85	92.31
	Correctness $\times$ 5 Majority	29.40	14.53	90.77
	Correctness $\times$ 5 Unanimous	41.20	15.82	81.52
	Correctness $\cup$ 5 Unanimous	54.32	23.08	56.25
	3-Iter Pipeline	73.20	20.00	16.00
o3-mini	Vanilla Prompt (Baseline)	24.00	14.29	96.92
	Correctness	28.60	15.24	98.46
	Correctness $\times$ 5 Majority	29.20	15.18	96.92
	Correctness $\times$ 5 Unanimous	33.00	15.56	93.85
	Correctness $\cup$ 5 Unanimous	30.00	15.16	95.38
	3-Iter Pipeline	34.40	15.84	93.85
o3	Vanilla Prompt (Baseline)	58.12	20.73	78.46
	Correctness	70.60	22.00	50.00
	Correctness $\times$ 5 Majority	73.15	25.87	56.92
	Correctness $\times$ 5 Unanimous	83.77	26.47	13.85
	Correctness $\cup$ 5 Unanimous	78.60	28.57	43.08
	1-Iter Pipeline	75.40	24.00	42.00
	3-Iter Pipeline	81.65	30.99	34.38
	5-Iter Pipeline	81.50	26.23	25.40
Multi-model ensemble	Correctness (5 Models) Majority	45.00	17.99	90.77
	Correctness (5 Models) Unanimous	78.60	25.00	32.31
	3-Iter Pipeline (2 Models) Unanimous	85.40	40.00	24.62

Finding 2: Attaining High Precision is Difficult

Metric	Answer Models			
	o3	Claude Sonnet 3.7	Gemini 2.5 Pro	GPT-4o
% answers passed UQ-Validator	0%	0%	12%	0%
% answers passed human reviewers (i.e., GT accuracy)	0%	0%	4%	0%
Human/UQ-Validator judgment agreement	100%	100%	92%	100%
Human-rated accuracy of UQ-Validator reasoning trace	96%	96%	76%	100%

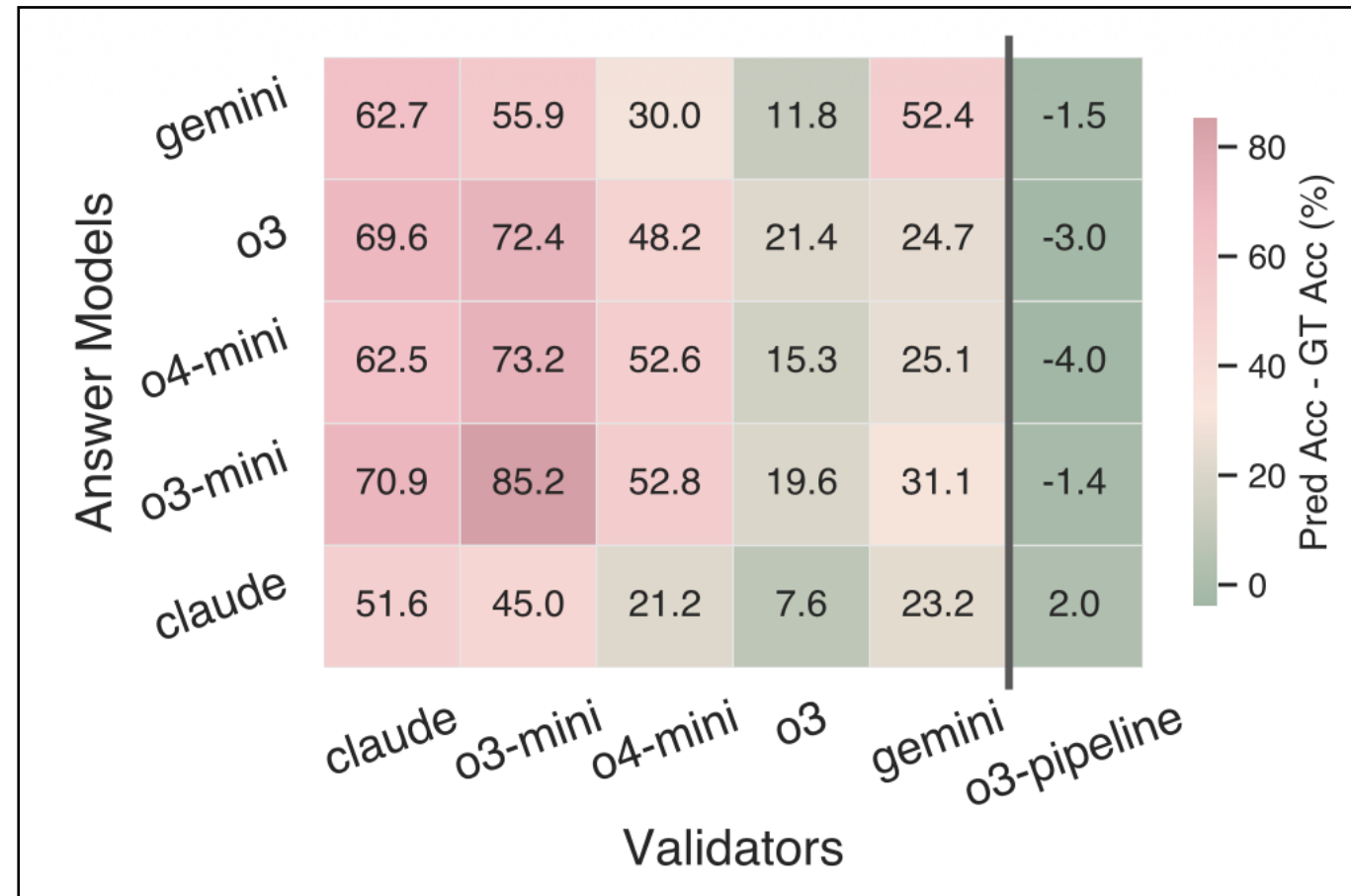
Nevertheless, the best resulting UQ-Validator remains useful for human reviewers!

Finding 3: Simple Validators Show Over-Optimism and Self-Bias



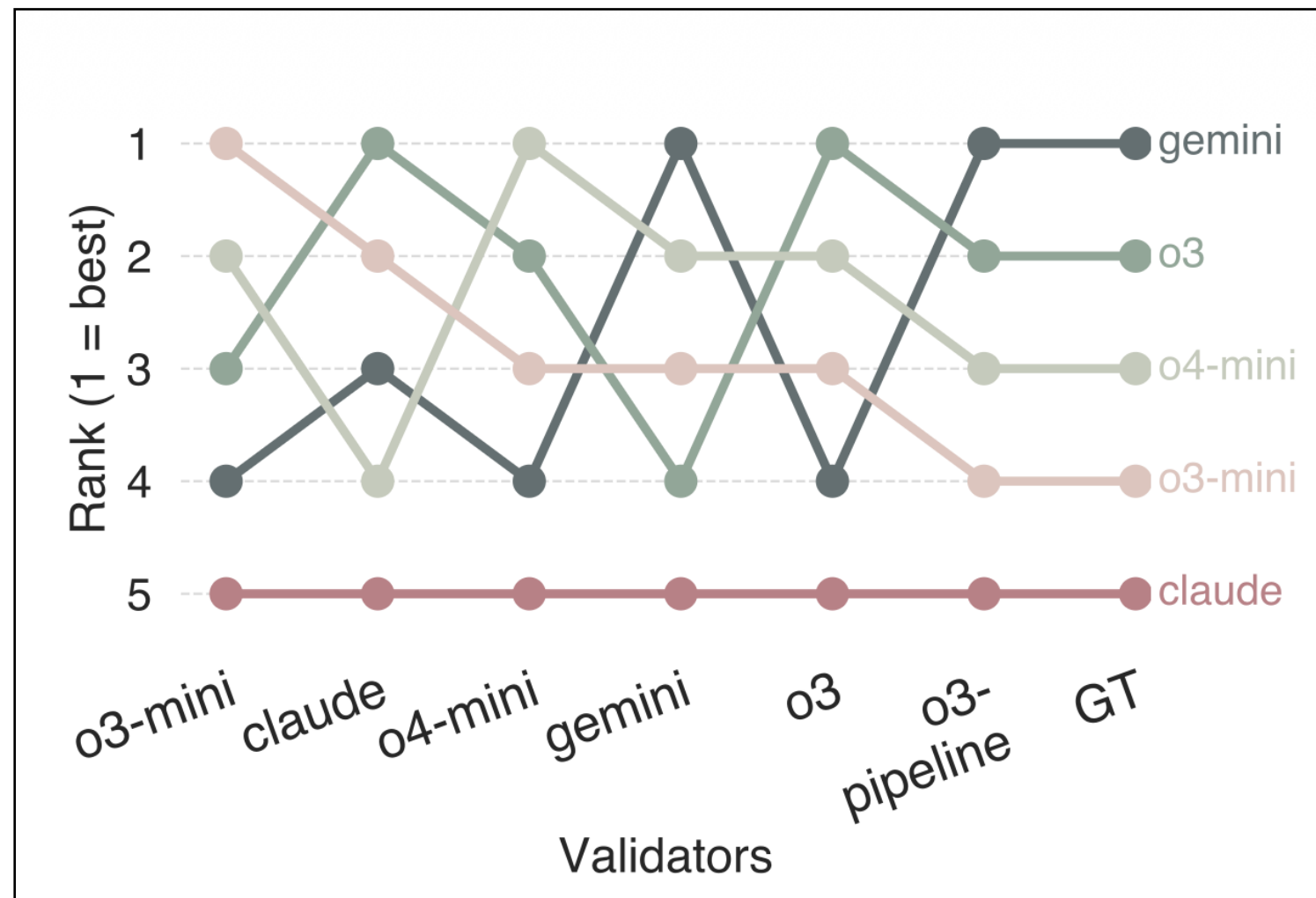
Simple LLM validators overrate self and sibling answers.

Finding 4: Compound Validator Strategies Mitigate Over-Optimism and Self-Bias



3-iter o3 pipeline largely removes over-optimism across all models!

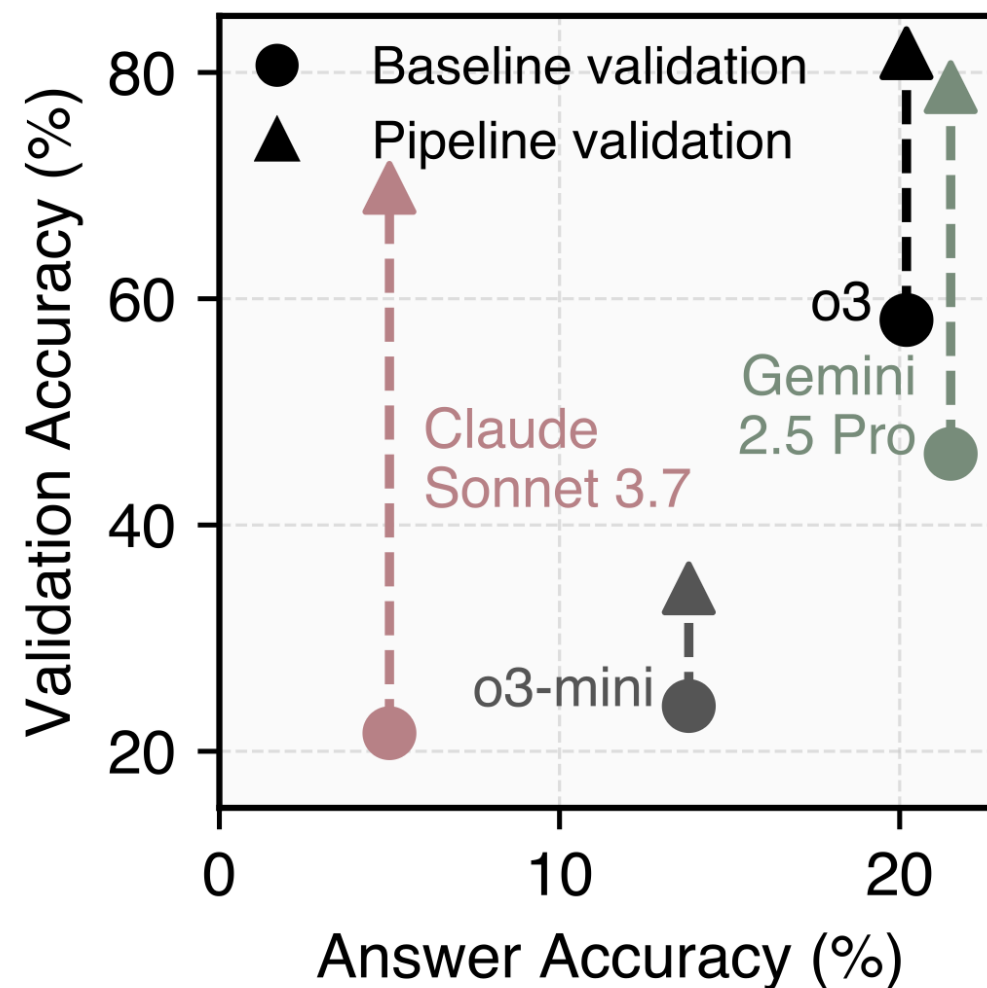
Finding 5: Model Rankings Are Unstable Across Validator Performance



Model ranking is unstable across validator performance

Using 3-iter pipeline aligns better with GT ranking

Finding 6: Better Answer Generators May Not Be Better Answer Validators



- Better answer performance is broadly indicative of better validation performance, but it's not always the case.

Oracle-free validator alone is not enough for unsolved questions.

Adding New Knowledge Requires Human Consensus!

UQ-Platform: An Open Platform for Community-Based Evaluation

UQ: Assessing Language Models on Unsolved Questions

[→ Login](#)[Paper](#)[Code](#)[Hugging Face Dataset](#)[Home](#)[Questions](#)[Contribute](#)[Contact](#)[Cite](#)

About Unsolved Questions

Benchmarks shape progress in AI research. A useful benchmark should be both *difficult* and *realistic*: questions should challenge frontier model while also reflecting real-world usage. Yet, current paradigms face a difficulty–realism tension: exam-style benchmarks are often made artificially difficult with limited real-world value, while benchmarks based on real user interaction often skew toward easy, high-frequency problems.

This work explores a radically different paradigm: assessing models on *unsolved* questions. Rather than a static benchmark scored once, we curate unsolved questions and evaluate models asynchronously over time with validator-assisted screening and community verification. We introduce **UQ**, a testbed of 500 challenging, diverse questions sourced from Stack Exchange, spanning topics from CS theory and math to less explored areas like sci-fi and history, probing capabilities including reasoning, factuality, and browsing. **UQ** is difficult and realistic by construction: unsolved questions are often hard and naturally arise when humans seek answers, thus solving them yields direct real-world value.

[Submit Model Answers](#)[All: 500 questions](#)[Technology: 52 questions](#)[Culture & Recreation: 16 questions](#)[Life & Arts: 35 questions](#)[Science: 395 questions](#)

Partial Model Evaluation

Answer Model	UQ-Validator Pass Rate		Human Pass Rate (2025-08-24)*
	# Passed	%	
O3-PRO	75 / 500	15.0%	4 / 46
↪ on UQ diamond subset	3 / 25	4.0%	0 / 2
GEMINI 2.5 PRO	25 / 500	5.0%	3 / 10
O4-MINI (HIGH)	25 / 500	5.0%	2 / 14
O3	44 / 500	8.8%	1 / 25
↪ on UQ diamond subset	1 / 25	2.0%	0 / 1
DEEPSEEK-R1-0528	11 / 500	2.2%	1 / 5
CLAUDE OPUS 4	7 / 500	1.4%	0 / 3
CLAUDE SONNET 3.7 (16K)	6 / 500	1.2%	0 / 3
GPT-4o	0 / 500	0.0%	0 / 0
Total unique questions	144 / 500	28.8%	10 / 91

Takeaways

- UQ introduces a new evaluation paradigm
 - Moving from create-once & run-once benchmarks to discovering new knowledge through continuous evaluation
- UQ has three components:
 - UQ-Dataset: 500 unsolved questions sourced from Stack Exchange
 - UQ-Validator: compound validation strategies to provide evaluation signals and pre-screen candidate solutions for human review
 - UQ-Platform: an open platform where experts collectively verify questions and solutions, enabling ongoing, asynchronous, and community-driven evaluation

Teams

UQ: Assessing Language Models on Unsolved Questions

Fan Nie^{♠*} Ken Ziyu Liu^{♠*}
Zihao Wang[♠] Rui Sun[♠] Wei Liu[♠]
Weijia Shi[♡] Huaxiu Yao[♣] Linjun Zhang[◇] Andrew Y. Ng[♠]
James Zou[♠] Sanmi Koyejo[♠] Yejin Choi[♠] Percy Liang[♠] Niklas Muennighoff^{♠*}
[♠]Stanford University [♡]University of Washington [♣]UNC Chapel Hill [◇]Rutgers University

Ongoing next step

- One-sided automated verification for open problems
- Self-improvement on unsolved questions with weak signals
-

Questions