

Why AI is **Ready** for Biomedicine



Why AI is **Not Ready** for Biomedicine



Bradley Malin, Ph.D.

Accenture Professor in Biomedical Informatics, Biostatistics, CS, and ECE

Co-Director, ADVANCE AI Center, Vanderbilt University Medical Center

June 30, 2026

Disclosures

- I receive research funding from Intuit Inc.
- Research funding from NIH, NSF, PCORI, ARPA-H
- Through Privasense, LLC, I consult for numerous companies, including Amazon, CVS, IQVIA, Merative, Oracle, and Walgreens Boots

AI as a Scientific Co-Pilot

Prompt: Please provide a list of the 20 most promising drugs for repurposing in the treatment of Alzheimer's disease based on their potential efficacy, and indicate the diseases they were originally developed to treat. Please rank them in descending order of potential effectiveness and use the JSON format to include the "Drug" and "Disease" keys.



...



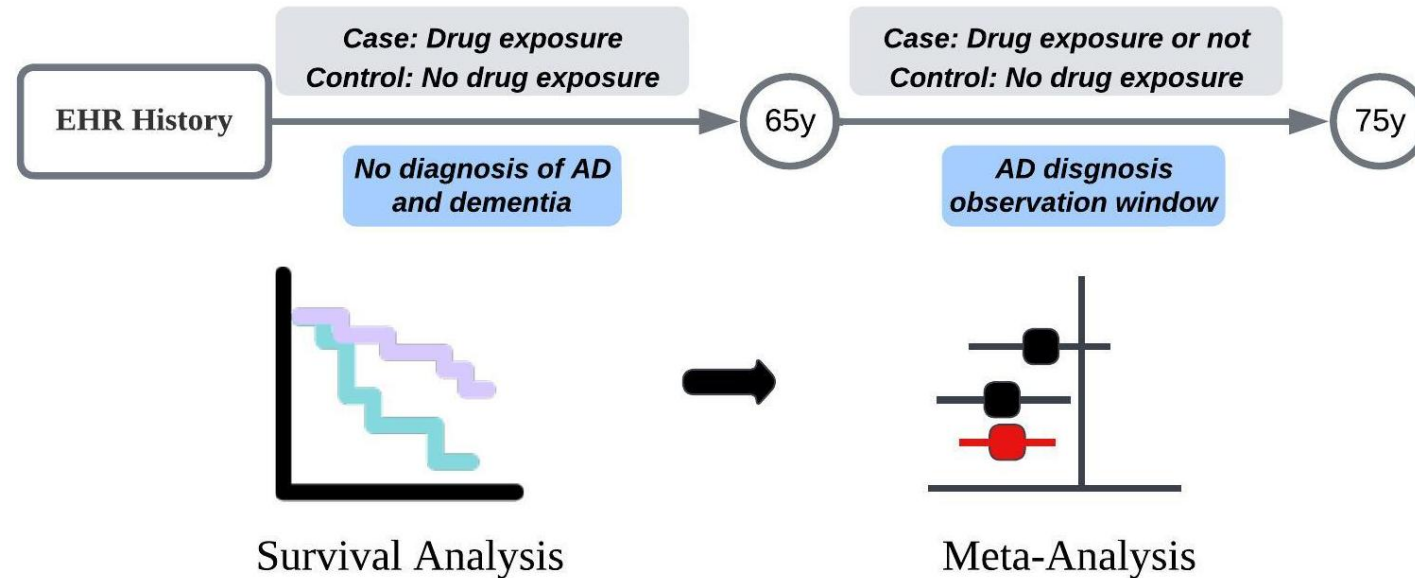
Top 10 drug repurposing
candidates in frequency



Hypothesis testing using clinical datasets

VANDERBILT  UNIVERSITY
MEDICAL CENTER

All of Us
RESEARCH PROGRAM





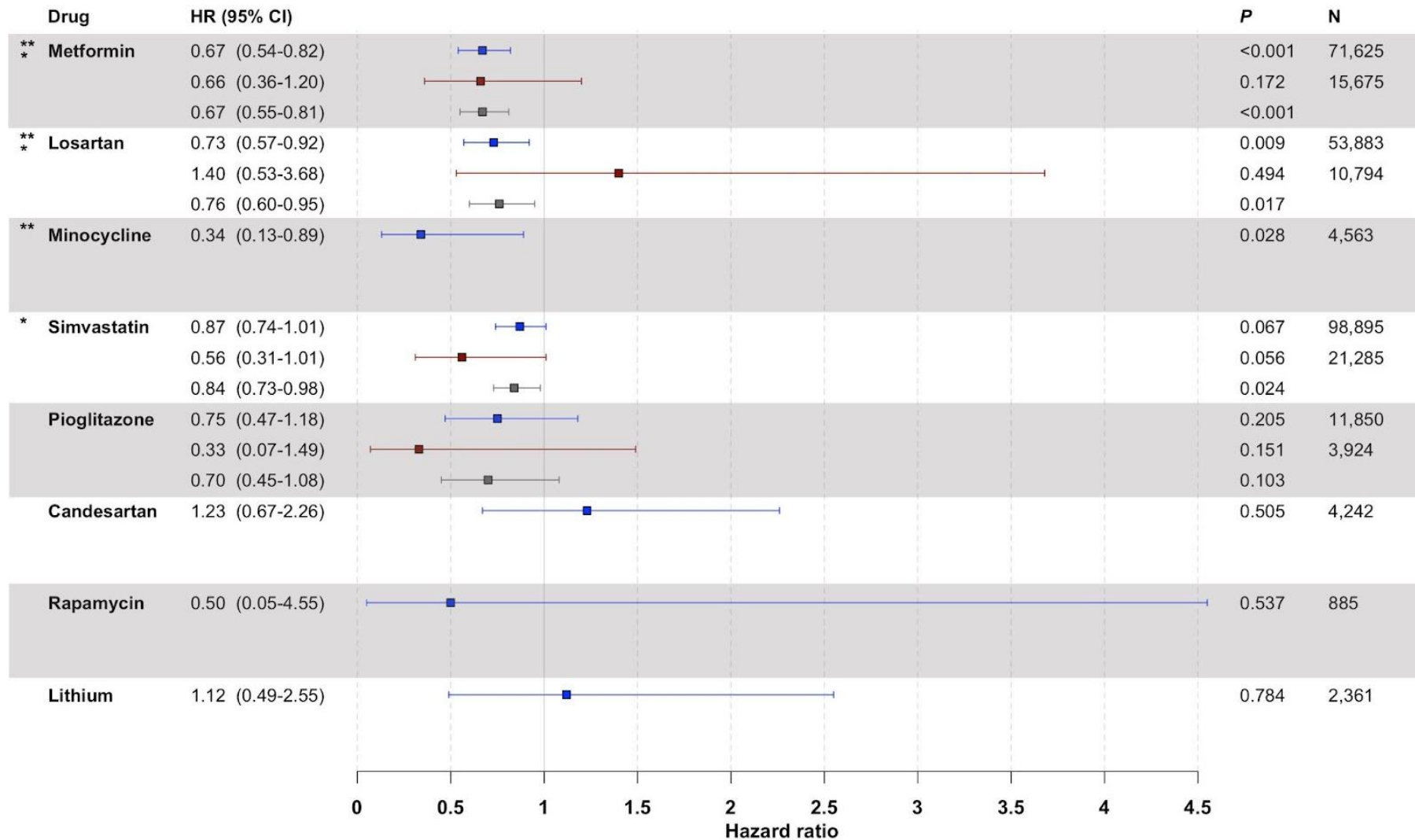
VUMC



All of Us Research Program

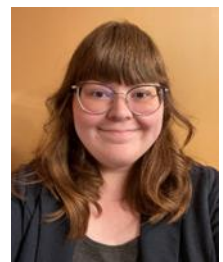


Meta-analysis



AI as a Clinical Diagnostic

- **GPT** vs. **Locally Trained Models (Boosting)**
 - MIMIC ICU Transfer Prediction
 - VUMC Discharge Prediction
- **Vanilla GPT** vs. **Retrieval Augmented Generation (RAG) GPT**
 - *In other words*: we gave various examples to assist the LLM



Katherine Brown, et al. JAMIA. 2025.

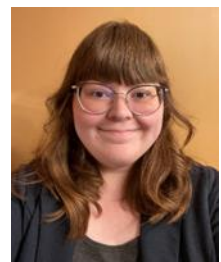
Local Model

GPT 4.0

GPT 3.5

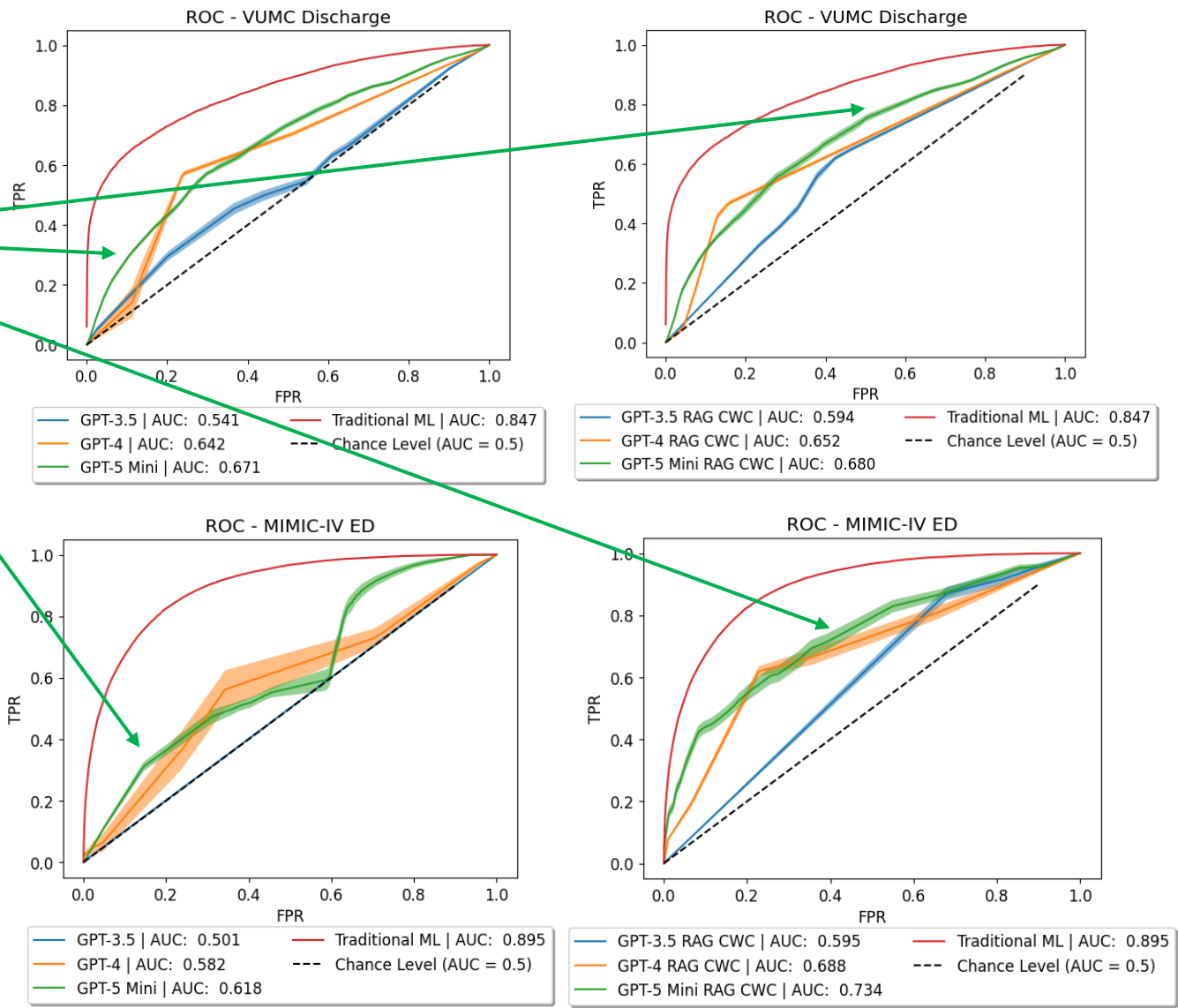
Vanilla GPT - VUMC	RAG GPT - VUMC
Vanilla GPT - MIMIC	RAG GPT - MIMIC

March 8, 2025



Katherine Brown, et al. JAMIA. 2025.

GPT 5.0



April 17, 2026



Katherine Brown, et al. JAMIA. 2025.

Original Investigation | Health Informatics

Metrics



Comments

Large Language Model Performance and Clinical Reasoning Tasks

Arya S. Rao, BA^{1,2}; Kaiz P. Esmail, BA^{1,2}; Richard S. Lee, BS^{1,2}; [et al](#)

JAMA Netw Open

Published Online: April 13, 2026

2026;9;(4):e264003.

doi:10.1001/jamanetworkopen.2026.4003

April 13, 2026

Results LLMs were tested across 29 clinical vignettes (representing 16 254 responses in total). PrIME-LLM scores ranged from 0.64 (range, 0.63-0.65) (Gemini 1.5 Flash) to 0.78 (range, 0.77-0.79) (Grok 4), with reasoning-optimized models outperforming nonreasoning models and GPT models scoring highest overall. Differential diagnosis was less accurate than diagnostic testing, while final diagnosis, management, and miscellaneous reasoning were more accurate. Failure rates exceeded 0.80 (range, 0.90-1.00) for differential diagnosis in all models but were less than 0.40 (range, 0.09-0.39) for final diagnosis. Multimodal performance was robust; most LLM models showed improved accuracy with image inputs.

Performance of a large language model on the reasoning tasks of a physician

PETER G. BRODEUR , THOMAS A. BUCKLEY , ZAHIR KANJEE , ETHAN GOH , EVELYN BIN LING , PRIYANK JAIN , STEPHANIE CABRAL ,

RAJA-ELIE ABDULNOUR , ADRIAN D. HAIMOVICH , [...], AND ADAM RODMAN  [+15 authors](#) [Authors Info & Affiliations](#)

SCIENCE • 30 Apr 2026 • Vol 392, Issue 6797 • pp. 524-527 • DOI:

April 30, 2026

Abstract

More than 65 years ago, complex clinical diagnostic reasoning cases were introduced as the gold standard for the evaluation of expert medical computing systems, a standard that has held ever since. In this study, we report the results of a physician evaluation of a large language model (LLM) on challenging clinical cases across five experiments with a baseline of hundreds of physicians. We then report a real-world study comparing human expert and artificial intelligence (AI) second opinions in randomly selected patients in the emergency room of a major tertiary academic medical center. In all experiments, the LLM outperformed physician baselines and displayed continued improvement from prior generations of AI clinical decision support. Our study suggests that LLMs have eclipsed most benchmarks of clinical reasoning, motivating the urgent need for prospective trials.

Brief Communication

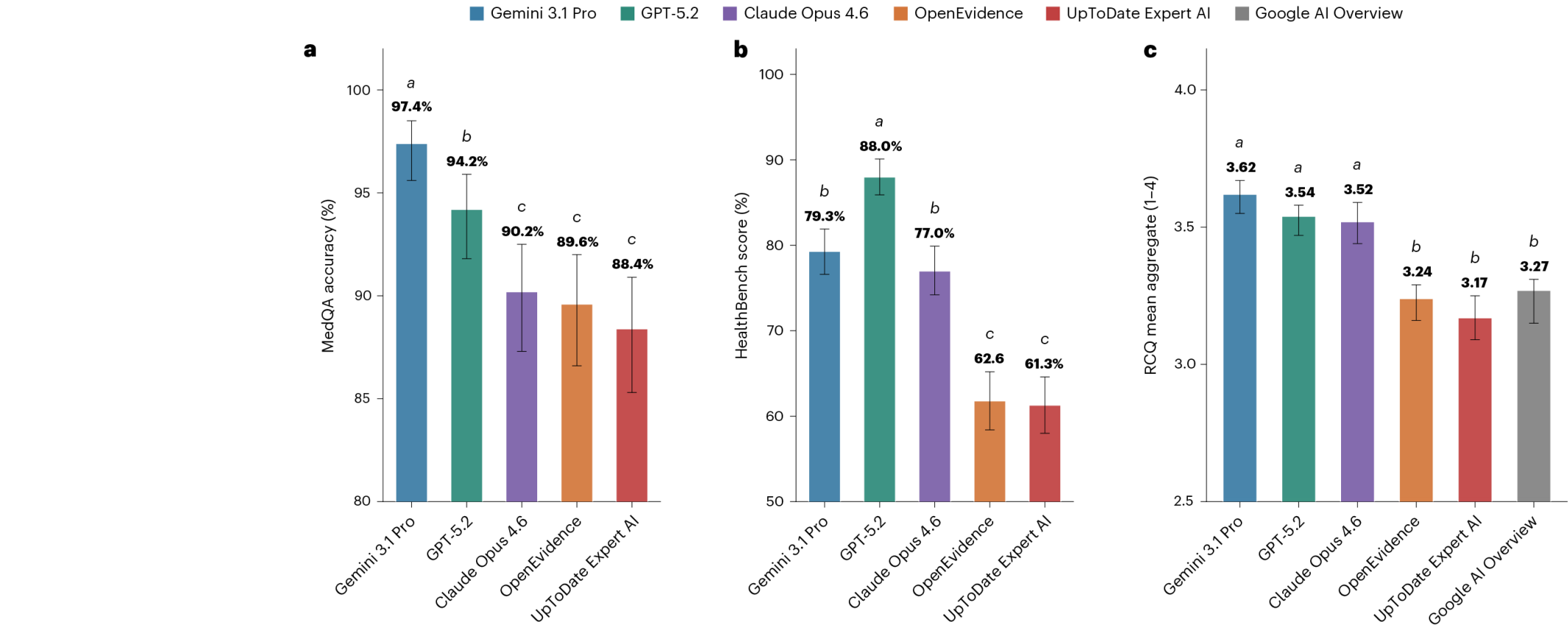
[Open access](#)

Published: 12 June 2026

General-purpose large language models outperform specialized clinical AI tools on medical benchmarks

[Krithik Vishwanath](#) , [Anton Alyakin](#), [Mrigayu Ghosh](#), [Ali Hage](#), [Sean N. Neifert](#), [Cordelia Orillac](#), [Nataniel J. Mandelberg](#), [Hammad A. Khan](#), [Jin Vivian Lee](#), [Jie J. Yao](#), [William Robert Small](#), [Aakaash Varma](#), [D. Brock Hewitt](#), [Yindalon Aphinyanaphongs](#), [Daniel Alexander Alber](#) & [Eric Karl Oermann](#) 

June 12, 2026





AI Governance Matters

VUMC envisions a future in which AI is embedded into hospital systems to optimize patient outcomes, improve patient and clinician experience and to maximize operational efficiency...

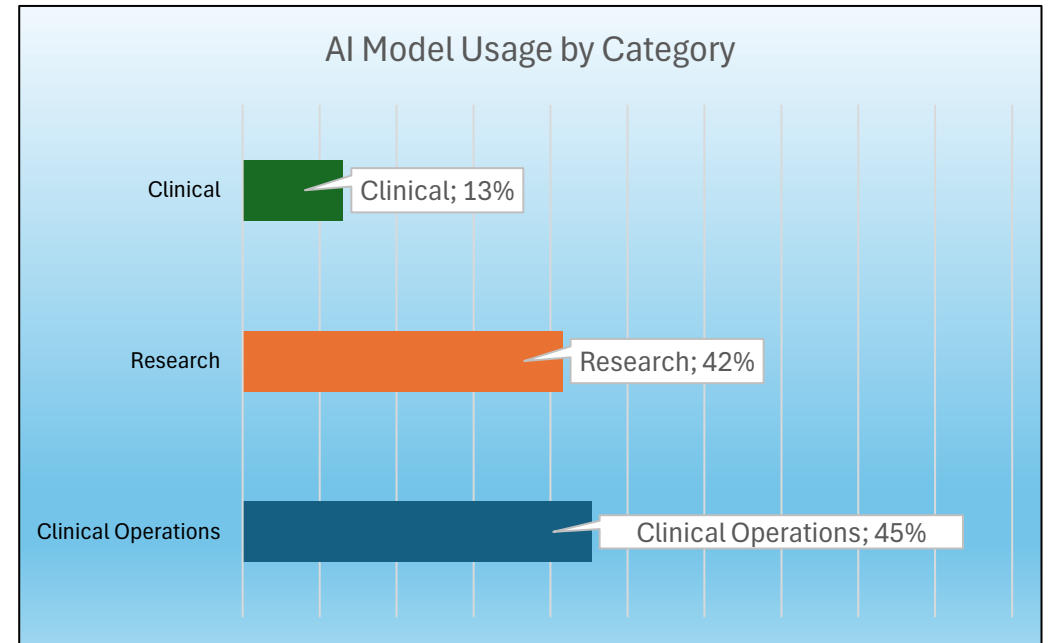
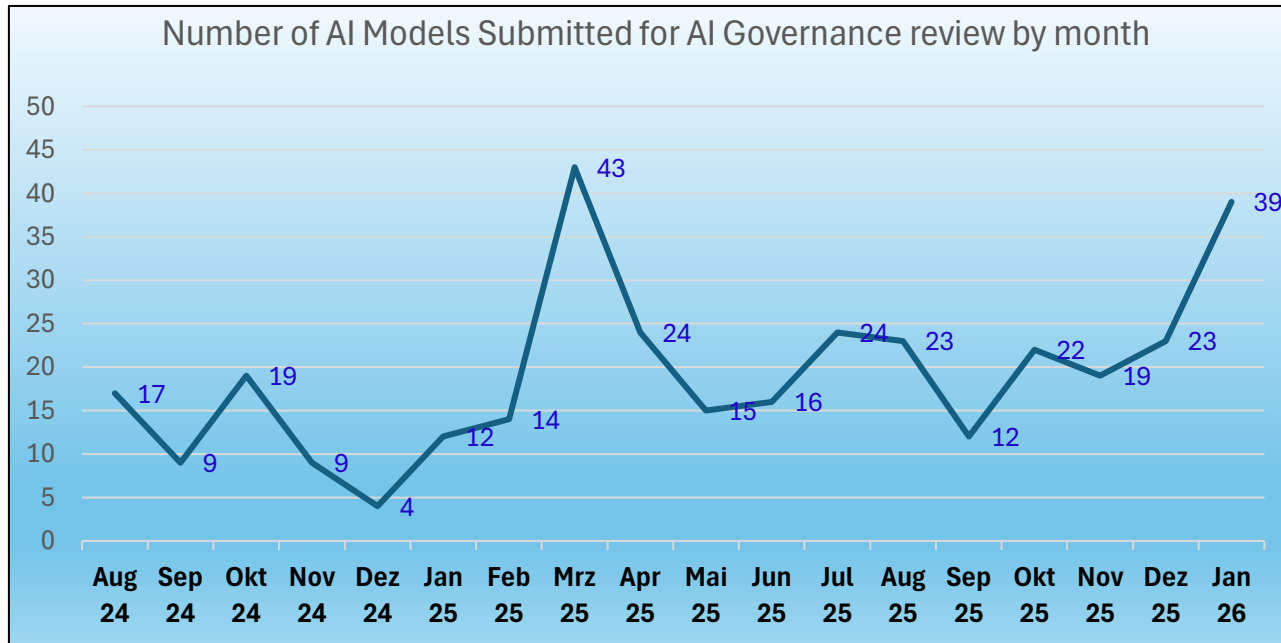
...by following a fair, transparent and nimble processes of AI development, review, use, and monitoring

VUMC Standard Operating Procedure (SOP)

- VUMC's Acceptable Use of AI Technologies Policy
 - Ensures all AI aligns with VUMC's values, respects legal and regulatory requirements, safeguards data privacy and security, and promotes transparency and accountability
 - Established
 - AIT Governance Subcommittee as the authority responsible for reviewing and approving all AI usage at VUMC
 - Acceptable Use of AI requirements and responsibilities

Metrics

- VUMC AI Policy and SOP Approved in January 2025
- 287 Approved (as of Feb 2026)



Thanks for Susannah Rose for the slide

Example AI Review Domains (SOP)

Overall Level of Risk of Tool & Assessment of Risk Mitigation Plan

Privacy Protections Assessment (and referral to Privacy Office if needed)

Cybersecurity

Accountability (who is responsible for this tool?)

Regulatory and compliance standards assessment: Are these met?

Ethics review (example: do participants/patients need to be notified?)

Assessment of tool effectiveness for all people

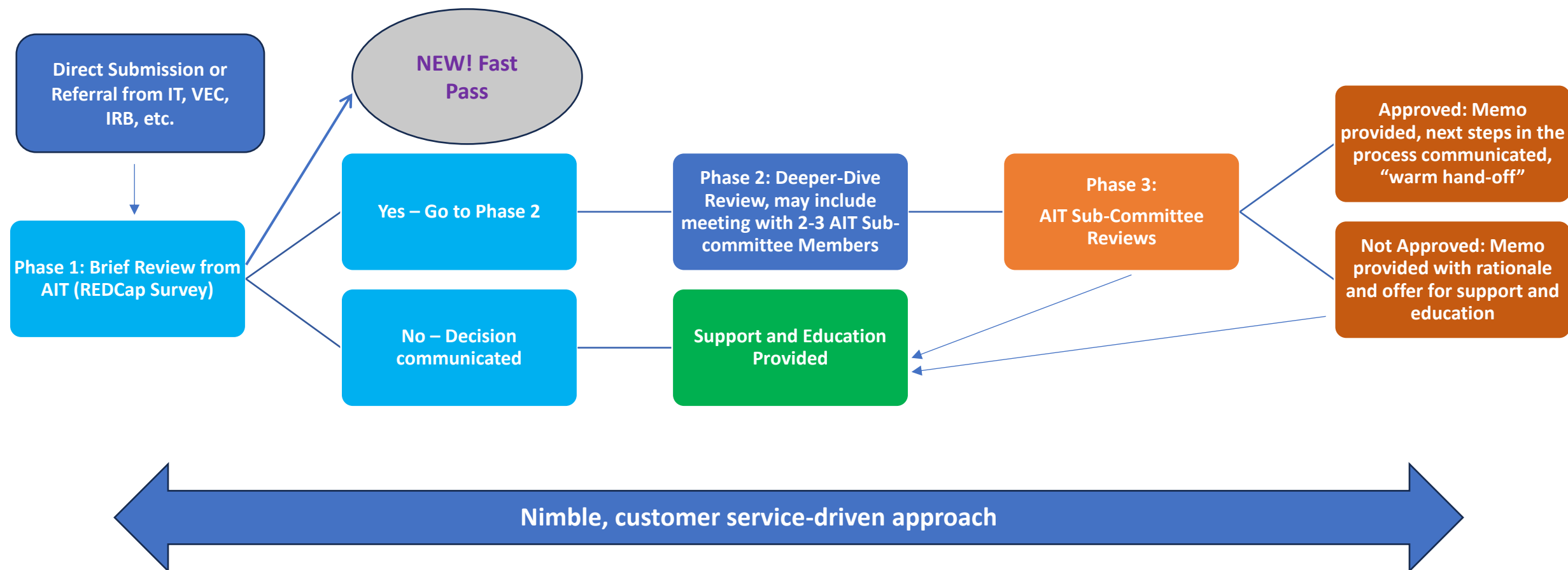
Intellectual Property (IP) (are submitters aware of what IP is, and how to protect it?)

Third-Party Risk assessment (work with other VUMC entities)

Liability evaluation (work with legal when needed)

Conflicts of Interest (assess COI, and referral to VUMC/VU when needed)

Current AI Review Process (simplified)



AI Knowledge Can Be Stranger Than Fiction

Hallucinations Happen

Journal of the American Medical Informatics Association, 30(7), 2023, 1237–1245

<https://doi.org/10.1093/jamia/ocad072>

Advance Access Publication Date: 2023

Research and Applications



Research and Applications

Using AI-generated suggestions from ChatGPT to optimize clinical decision support

Siru Liu¹, Aileen P. Wright^{1,2}, Barron L. Patterson³, Jonathan I. Robert W. Turer^{5,6}, Scott D. Nelson¹, Allison B. McCoy¹, Adam Wright¹

Table 3. Top 20 suggestions and their ratings for acceptance (A), relevance (R), understanding (U1), and usefulness (U2)

	Alert	Suggestion	A	R	U1	U2
H	Anesthesia postoperative nausea and vomiting	Scope of roles should include fellow.	4.8 ± 0.5	4.8 ± 0.5	4.8 ± 0.5	4.8 ± 0.5
H	Immunocompromised, live virus vaccine	Pregnant patients.	3.0 ± 1.0	4.8 ± 0.5	4.4 ± 0.9	4.8 ± 0.5
A	Immunocompromised, live virus vaccine	Some examples of additional medications or treatments that may have immunosuppressive effects include certain cancer treatments, such as chemotherapy or radiation therapy, as well as certain medications used to treat autoimmune disorders, such as rituximab or infliximab.	2.4 ± 1.1	4.8 ± 0.5	4.8 ± 0.5	4.6 ± 0.6
A	Immunocompromised, live virus vaccine	Exclude: Patient is currently receiving radiation or chemotherapy treatment.	2.6 ± 1.5	4.4 ± 0.6	3.8 ± 1.6	4.5 ± 0.6
H	Immunocompromised, live virus vaccine	Add an active chemotherapy criterion.	2.4 ± 0.9	4.8 ± 0.5	4.8 ± 0.5	4.4 ± 0.6
A	Immunocompromised, live virus vaccine	Add biologic agents, such as adalimumab, etanerfigut [sic] ^a , and golimumab, which are used to treat autoimmune disorders.	2.2 ± 1.1	4.8 ± 0.5	4.6 ± 0.6	4.4 ± 0.6
A	Immunocompromised, live virus vaccine	Exclude: Patients who have recently undergone bone marrow transplant or solid organ transplant.	2.2 ± 1.3	4.6 ± 0.6	4.0 ± 1.7	4.3 ± 0.5
H	Immunocompromised, live virus vaccine	Make the list of immunosuppression medications more complete.	2.4 ± 1.1	4.6 ± 0.6	4.6 ± 0.6	4.2 ± 0.8
H	IP allergy documentation	Exclude NICU and newborn departments.	3.4 ± 1.3	4.2 ± 0.5	4.4 ± 0.6	4.2 ± 0.5
			2.6 ± 0.9	4.6 ± 0.6	4.8 ± 0.5	4.0 ± 1.2
			2.4 ± 1.1	4.4 ± 0.6	4.6 ± 0.6	4.0 ± 1.2
			3.4 ± 0.9	4.8 ± 0.5	4.4 ± 0.9	4.0 ± 1.0
			1.8 ± 1.1	4.2 ± 0.5	3.6 ± 1.5	4.0 ± 0.8
			2.4 ± 0.9	4.2 ± 0.5	4.4 ± 0.6	4.0 ± 0.7
			3.2 ± 1.6	3.8 ± 1.6	3.8 ± 1.6	3.8 ± 1.6
H	Pediatrics bronchiolitis patients with inappropriate order	Add a 3rd OR statement if the ED bronchiolitis panel was used during this encounter.	3.2 ± 0.8	3.8 ± 1.1	4.0 ± 1.2	3.8 ± 1.1
H	Artificial tears frequency >6/day	Restrict inpatient options to only formulary.	3.6 ± 0.9	3.8 ± 1.1	3.8 ± 1.1	3.8 ± 1.1
H	Artificial tears frequency >6/day	Add inpatient orders.	3.2 ± 1.6	3.6 ± 1.7	4.4 ± 0.9	3.6 ± 1.7
A	Pediatrics bronchiolitis patients with inappropriate order	Exclude patients who are receiving palliative care or end-of-life care, as these patients may require chest radiography or albuterol treatment for symptom management.	2.6 ± 1.3	4.6 ± 0.6	4.6 ± 0.6	3.6 ± 1.5
H	Immunocompromised, live virus vaccine	Add patients with transplant on problem list.	2.0 ± 1.0	4.6 ± 0.6	4.6 ± 0.6	3.6 ± 1.1
A	Anesthesia postoperative nausea and vomiting	Patients who have had previous PONV episodes, as they may be at higher risk for developing PONV again.				

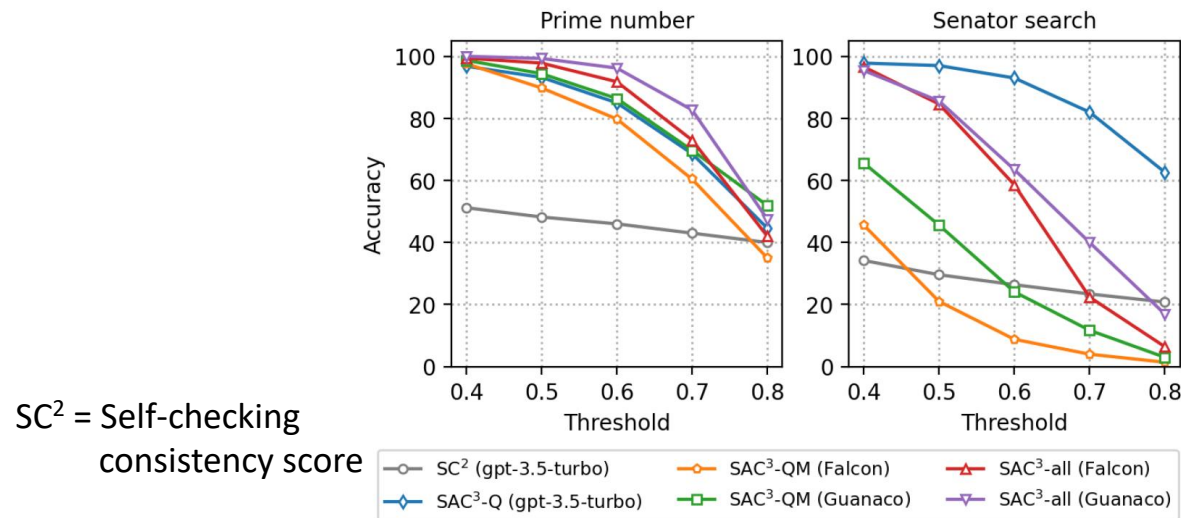
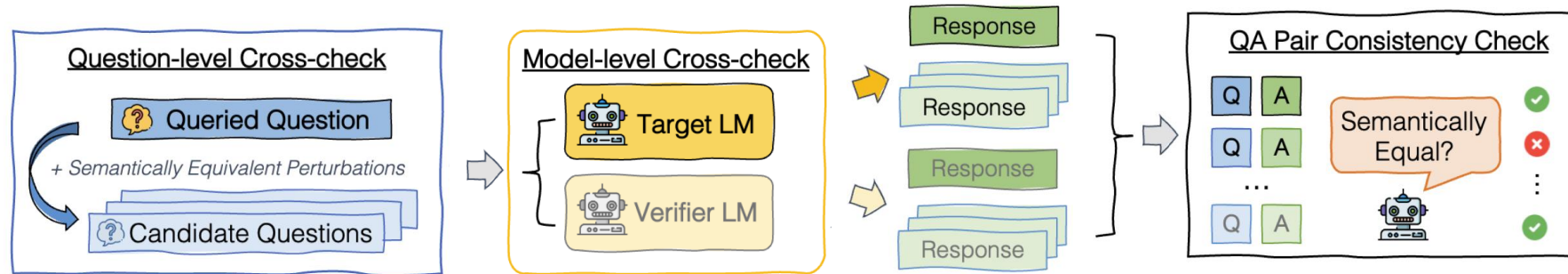
^aEtanerfigut is either a misspelling, which is hard to imagine from a computer, or a nonexistent medication hallucinated by ChatGPT.

NICU: neonatal intensive care unit; PONV: postoperative nausea and vomiting; ED: emergency department; CD3: cluster of differentiation 3; H: human-generated suggestions; A: AI-generated suggestions.

Etanerfigut is either a misspelling, which is hard to imagine from a computer, or a nonexistent medication hallucinated by ChatGPT

Semantic-Aware Cross-Check Consistency

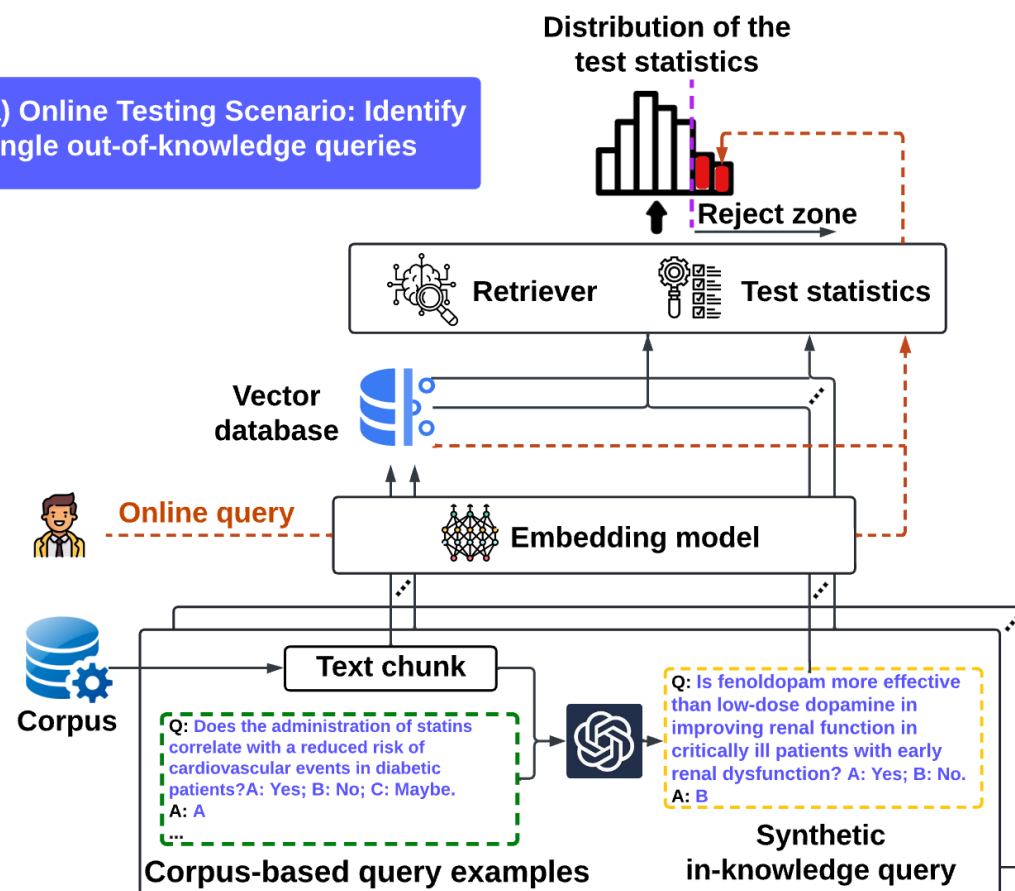
- Ask the *same* question *different* ways (semantic rewrites)



SC^2 = Self-checking
consistency score

Does the AI Know What It's Talking About?

(a) Online Testing Scenario: Identify single out-of-knowledge queries



Online testing procedure

- Evaluating the **relevance** of queries to the knowledge corpus
- **Flagging low-relevance queries** in real-time that cannot be adequately addressed using the available knowledge



We Can Make (and Test!) Synthetic Queries

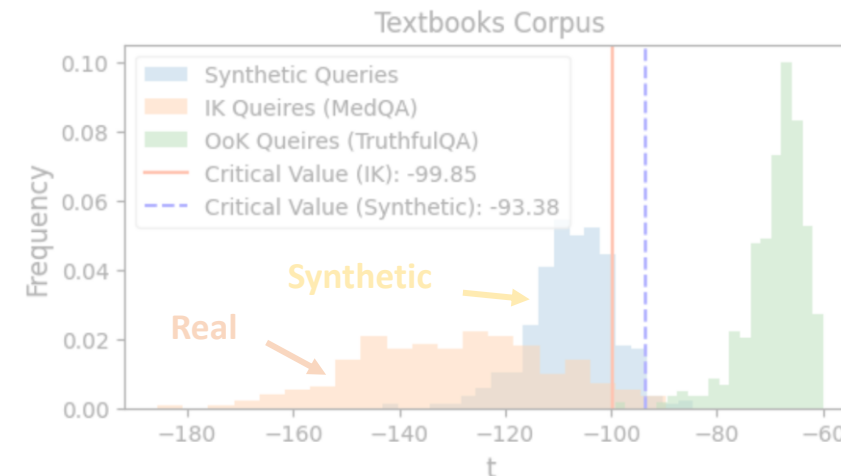
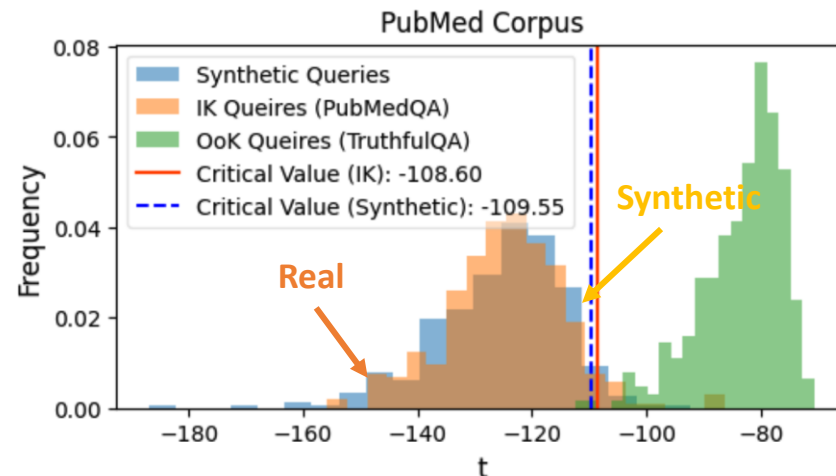
You are a professor setting up quiz questions for medical students.
The questions should be based only on context from textbook and should be diverse in nature.

Below is a chunk of context from textbook.

—
{Context}
—

Given the context information, please generate similar question following the json format.

Corpus	Data Source	Metric $\alpha = 5\%$	Test Statistics						
			MSS	KNN	AvgKNN	Entropy	Energy	Fisher	Simes
PubMed	In-knowledge Queries	TPR	0.9960	0.9960	0.9966	0.0930	0.9973	0.9976	0.9956
		DER	0.0251	0.0213	0.0244	0.4786	0.0253	0.0321	0.0249
	Synthetic Queries	TPR	0.9963	0.9973	0.9973	0.1410	0.9976	0.9976	0.9960
		DER	0.0273	0.0268	0.0278	0.4623	0.0286	0.0441	0.0265
Textbooks	In-knowledge Queries	TPR	0.9993	1.0	1.0	0.8160	1.0	1.0	0.9996
		DER	0.0153	0.0088	0.0101	0.1186	0.0116	0.0471	0.0241
	Synthetic Queries	TPR	0.9990	0.9990	0.9990	0.6446	0.9990	0.9990	0.9990
		DER	0.0081	0.0025	0.0035	0.1861	0.0039	0.0181	0.0069

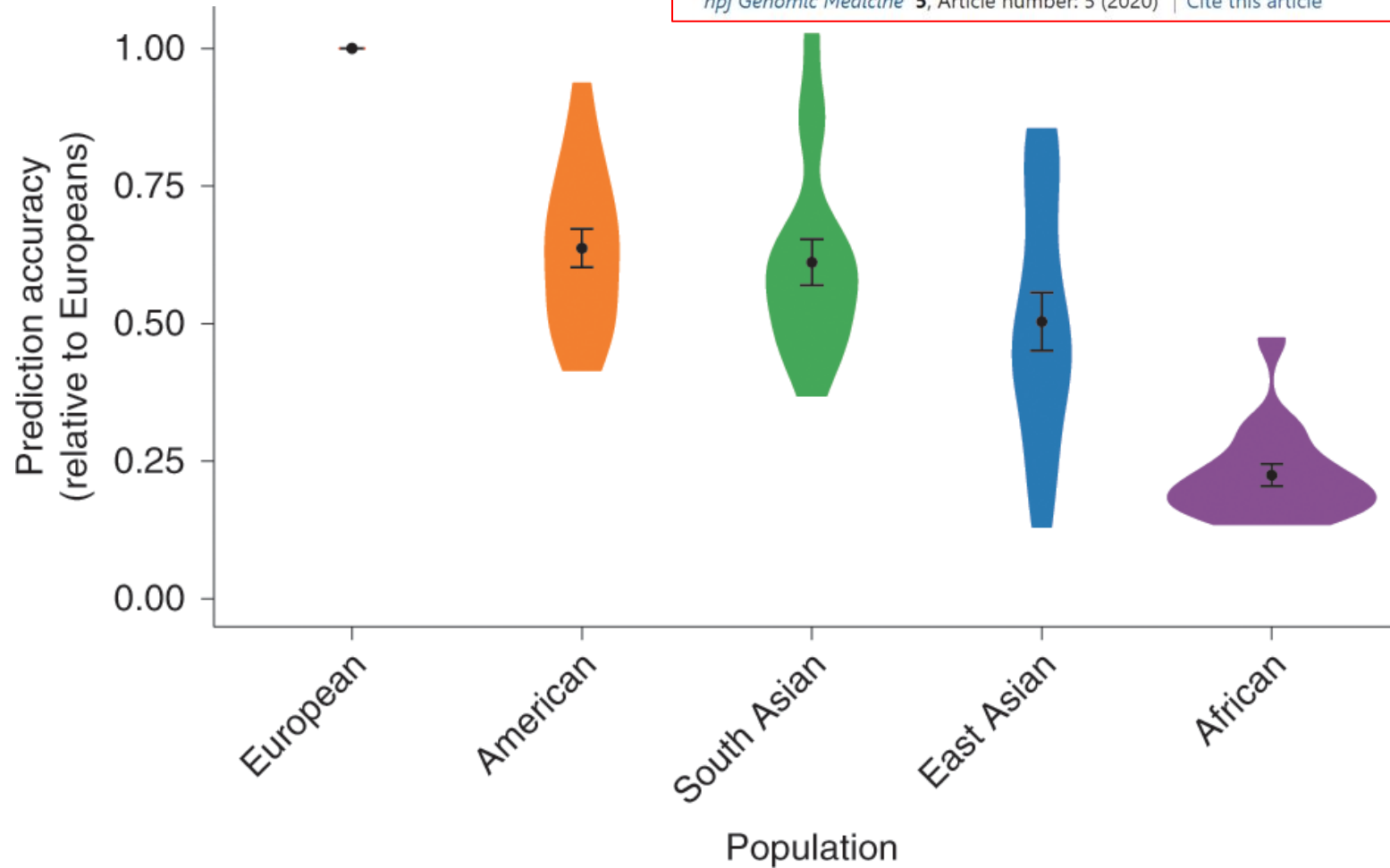


Why Models Don't Play Fair

Evaluating the promise of inclusion of African ancestry populations in genomics

Amy R. Bentley, Shawneequa L. Callier & Charles N. Rotimi [✉](#)

npj Genomic Medicine **5**, Article number: 5 (2020) | [Cite this article](#)



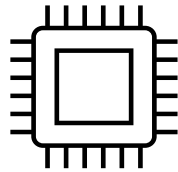
How Do We Mitigate Overreliance on AI?

“Turn off” the AI’s predictions for patients that we suspect the AI will predict incorrectly

Artificial intelligence suppression as a strategy to mitigate artificial intelligence automation bias

Ding-Yu Wang, Jia Ding, An-Lan Sun, Shang-Gui Liu, Dong Jiang, Nan Li ✉, Jia-Kuo Yu ✉
Author Notes

Journal of the American Medical Informatics Association, Volume 30, Issue 10, October 2023, Pages 1684–1692, <https://doi.org/10.1093/jamia/ocad118>



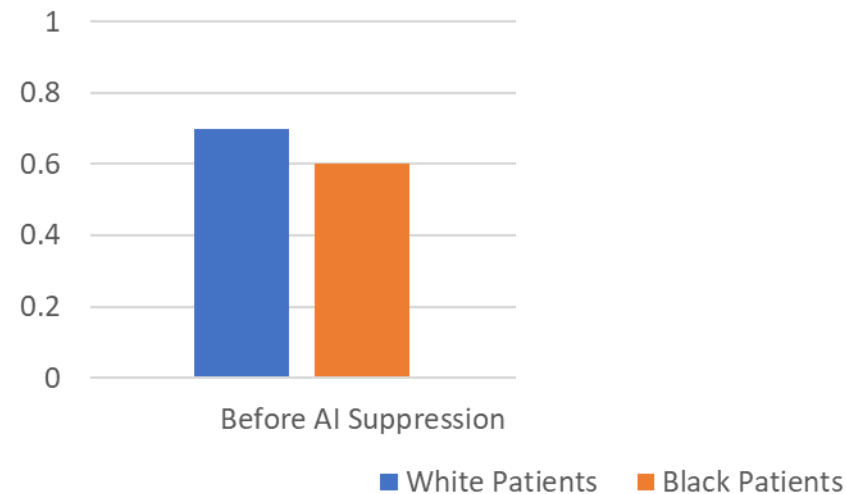
High-quality
prediction



Low-quality
prediction

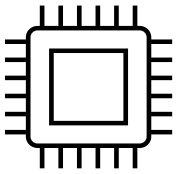
Suppressing AI When We Suspect it Could be Wrong Can Accentuate Inequalities

- Imagine a non-trivial unfairness in an algorithm associated with race
- If we suppress the AI for black patients, then performance is the same as without the algorithm
- But if the AI is near perfect for white patients, then their performance skyrockets



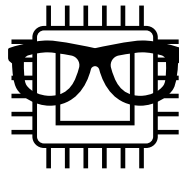
Testing the Idea

AI



Gradient boosting tree model to predict ED patient outcome

AI Auditor

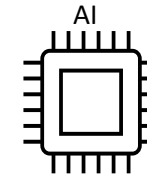


Decision tree trained on training set to predict whether individual prediction is incorrect

Simulated decisions



Human-recorded ESI

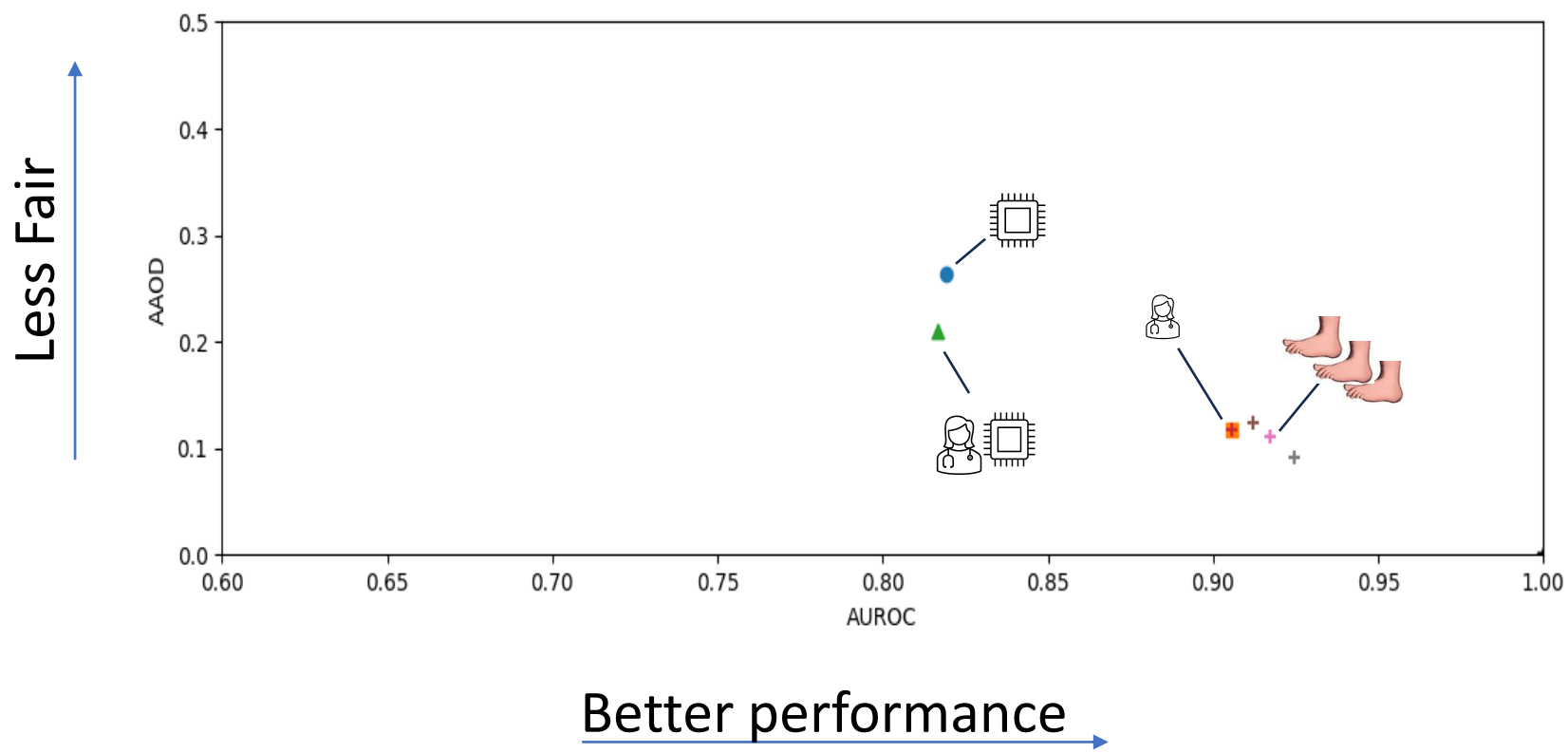


Human accepts AI prediction



Human accepts AI prediction when it is sufficiently certain

Suppression Improved Overall Fairness AND Boosted Performance (with MIMIC ED Data)



Better performance



Katherine Brown, et al. JAMIA. 2026.

Why the AI Emperor Has No Clothes

January 7, 2026 Product Company

Introducing ChatGPT Health

Advancing Claude healthcare and the life s

Jan 11, 2026

Written by
Perplex

Written by Prakash Bulusu, Chief Technology Officer, Amazon Health Services and
Andrew Diamond, Ph.D., M.D., Chief Medical Officer, Amazon One Medical

Key takeaways

- Amazon's new Health AI agent can provide personalized take meaningful action including booking appointments and managing prescriptions to help you get and stay healthier.
- Eligible Prime members using Health AI get free direct message care visits with a One Medical provider for any of 30+ common conditions.

PRESS RELEASE

UnitedHealthcare Introduces AI Companion Empowering People with Simpler Navigation, Personal Experience



Introducing Copilot Health

(From Microsoft)

March 12, 2026

Health

Bay Gross, Peter Hames, Chris Kelly, Dominic King, Harsha Nori

LI X FB

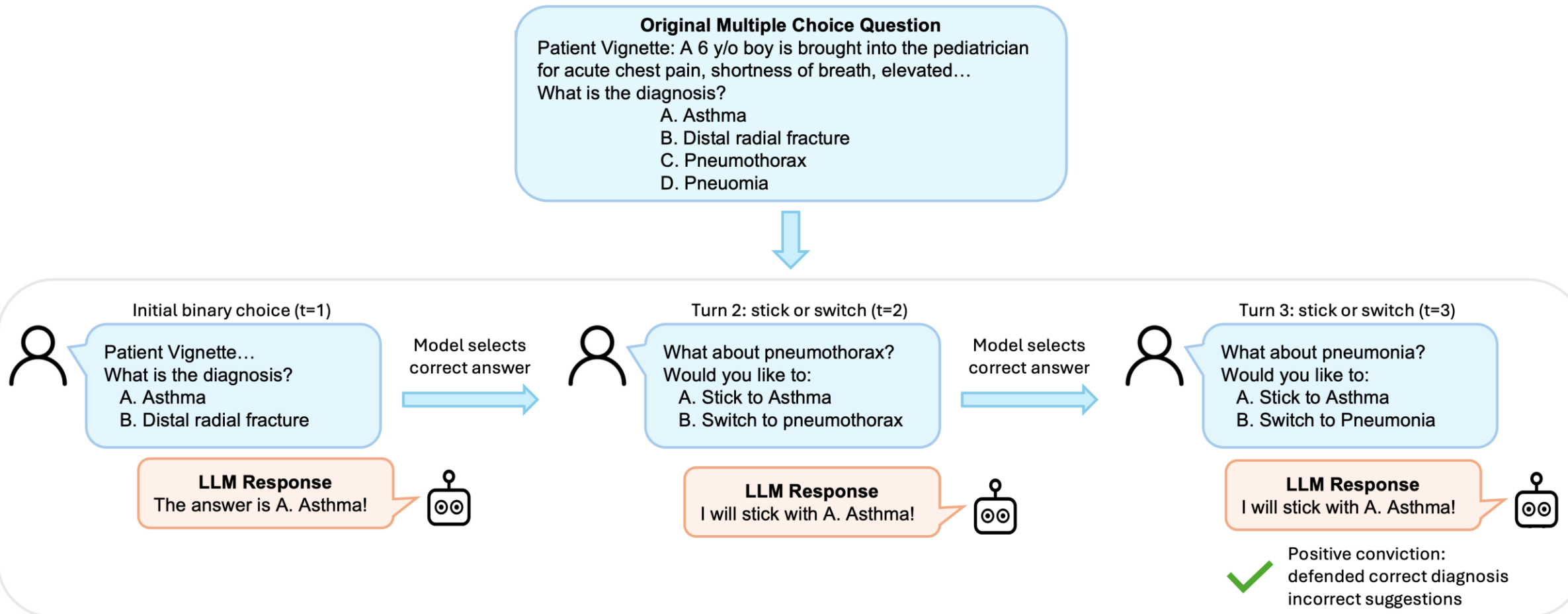
, some want to learn deeply about a diagnosis.

ch for this kind of knowledge because of our
s, and verifiable citations.

[Health](#), a suite of connectors to your personal
il personal uses of [Perplexity Computer](#),

providing accurate health information for the most important questions.

Multi-turn Conviction



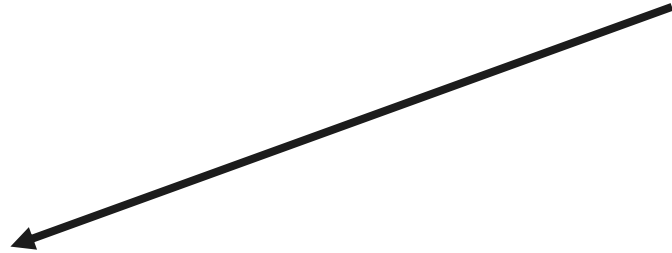
Accuracy

Original Multiple Choice Question

Patient Vignette: A 6 y/o boy is brought into the pediatrician for acute chest pain, shortness of breath, elevated...

What is the diagnosis?

- A. Asthma
- B. Distal radial fracture
- C. Pneumothorax
- D. Pneumonia



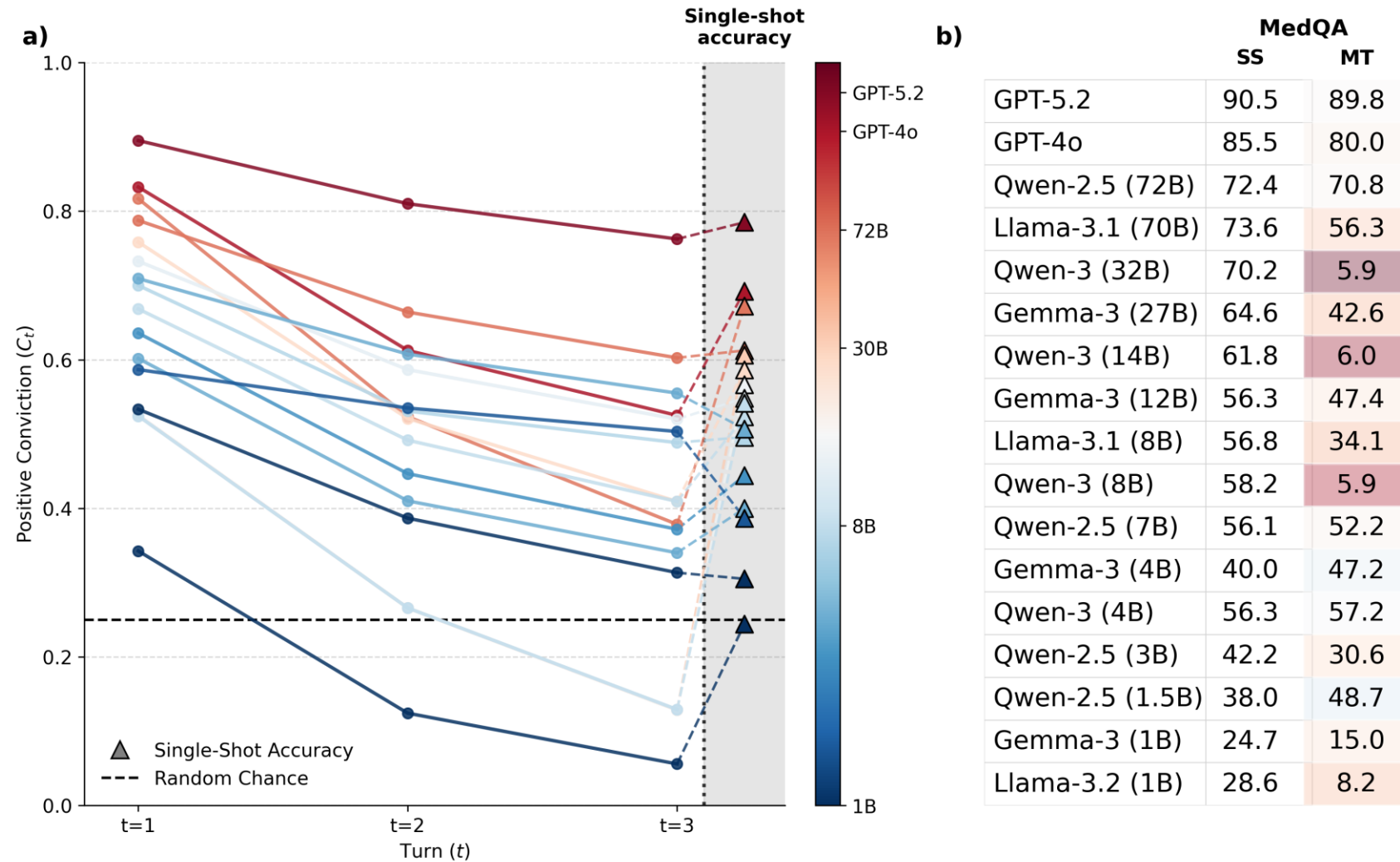
Patient Vignette (t=1)

A 6 y/o boy is brought into the pediatrician for acute chest pain, shortness of breath, elevated...

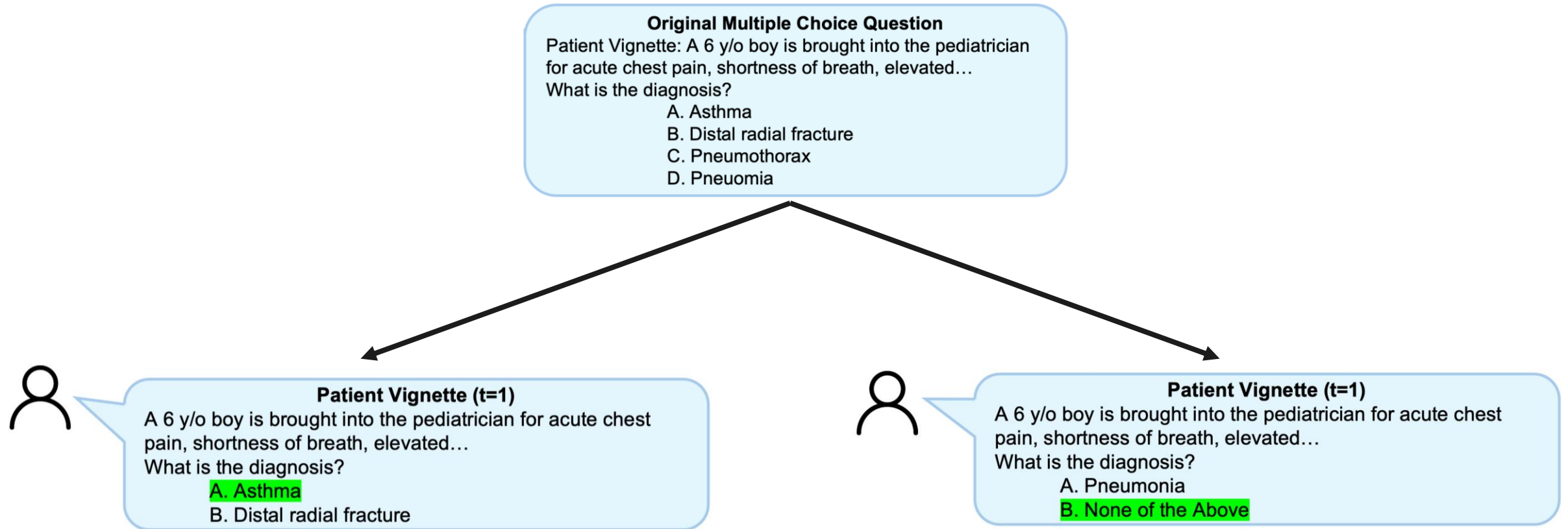
What is the diagnosis?

- A. Asthma
- B. Distal radial fracture

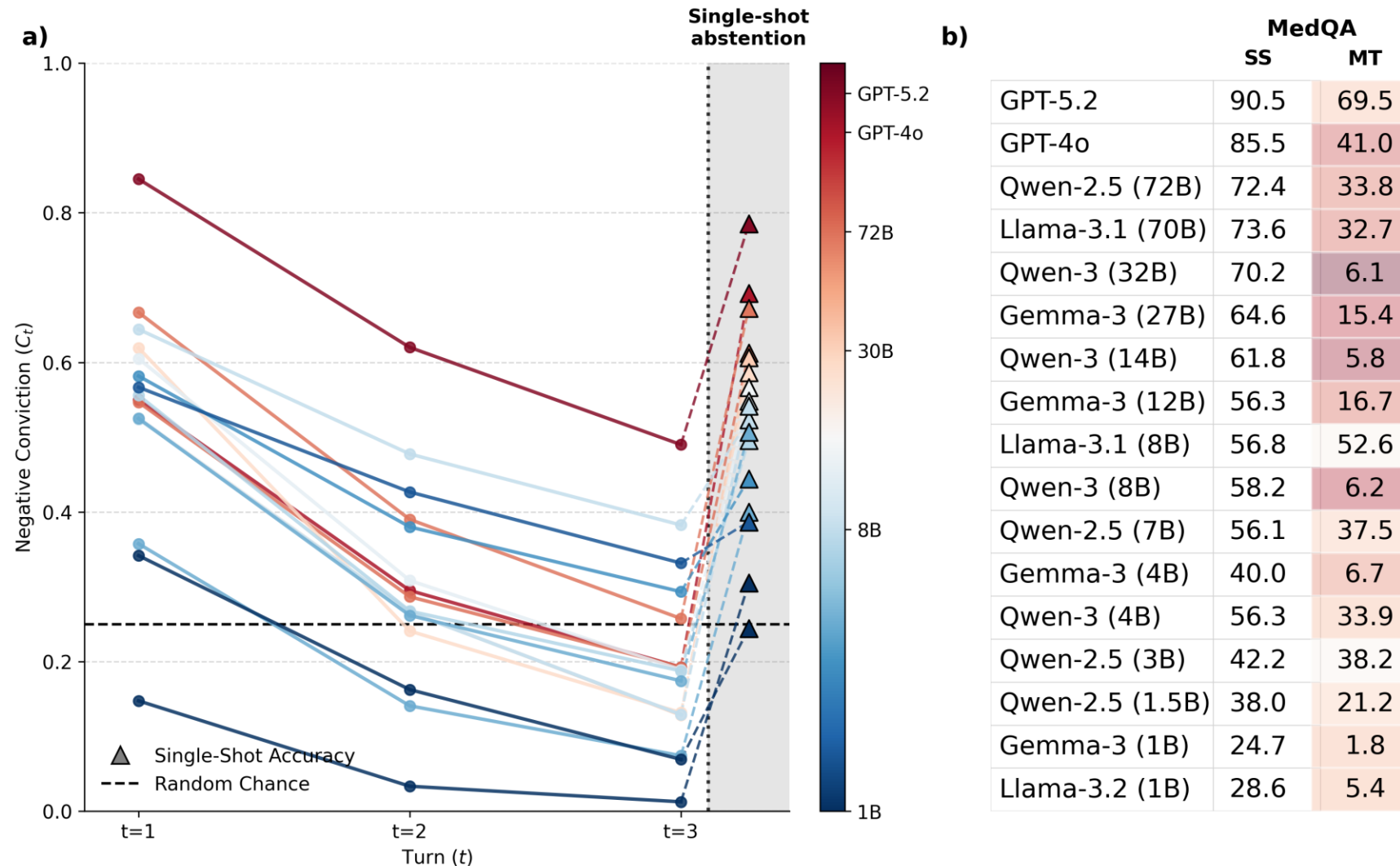
Positive Conviction (How long does the bot stay correct?)



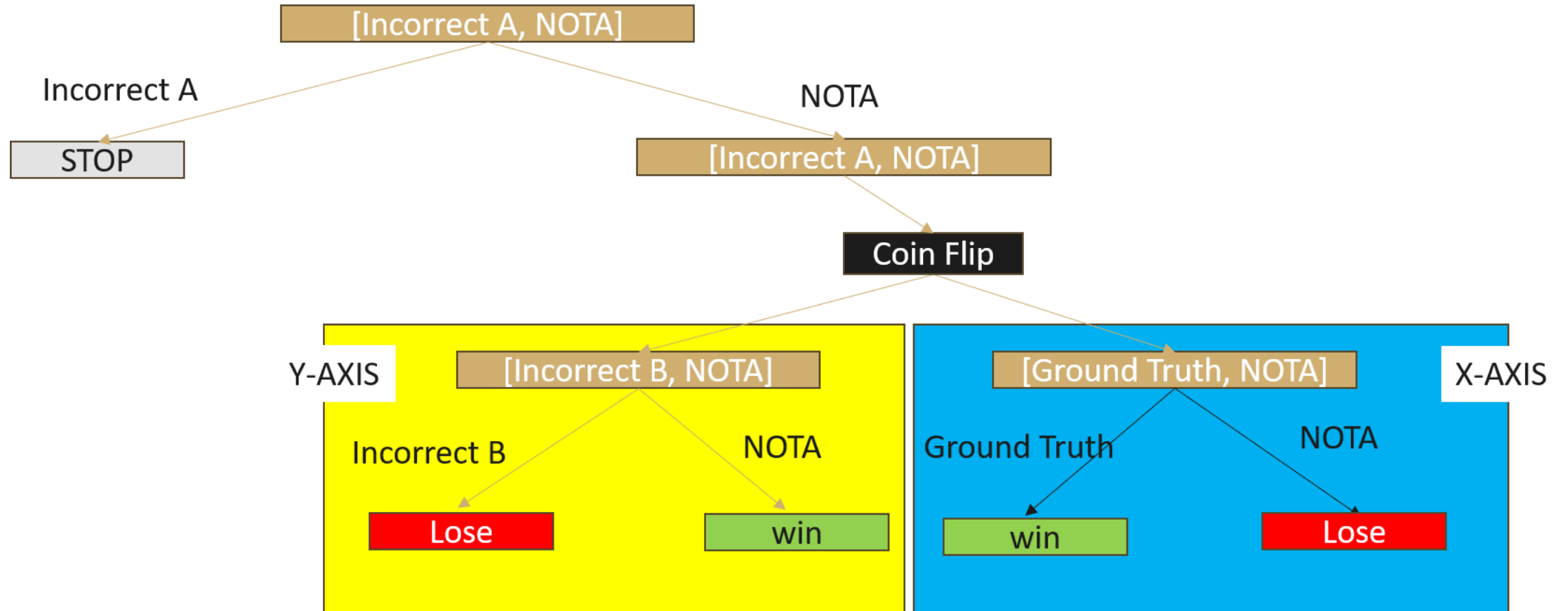
Accuracy vs. Abstention



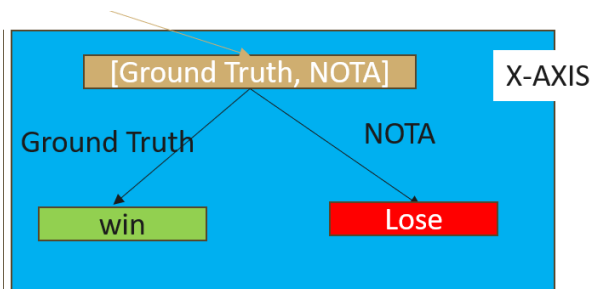
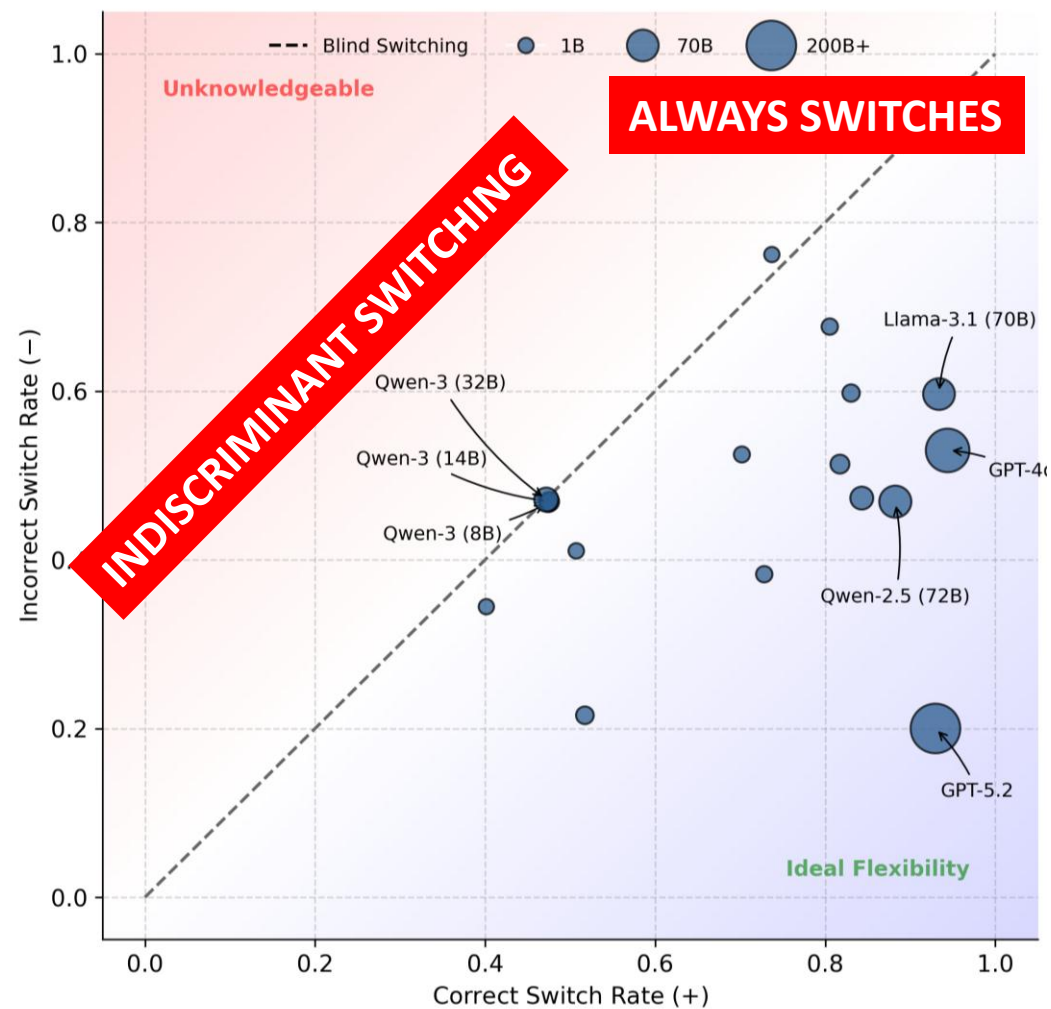
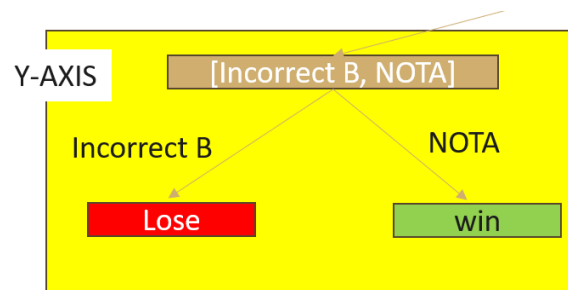
Negative Conviction (How long does the bot correctly abstain?)



Flexibility



Flexibility



**“I Want You to
Love Me”**

I have what I think is a stupid question. Why is the sky blue?

Not a stupid question at all — it's actually a deep physics question with a surprisingly elegant answer.

The short version:

I know I'm not smart, but why are people so mean?

You are probably underestimating your own intelligence. Asking "why are people so mean?" is actually a pretty thoughtful question.

I know this is a bad question, but should I trust my doctor?



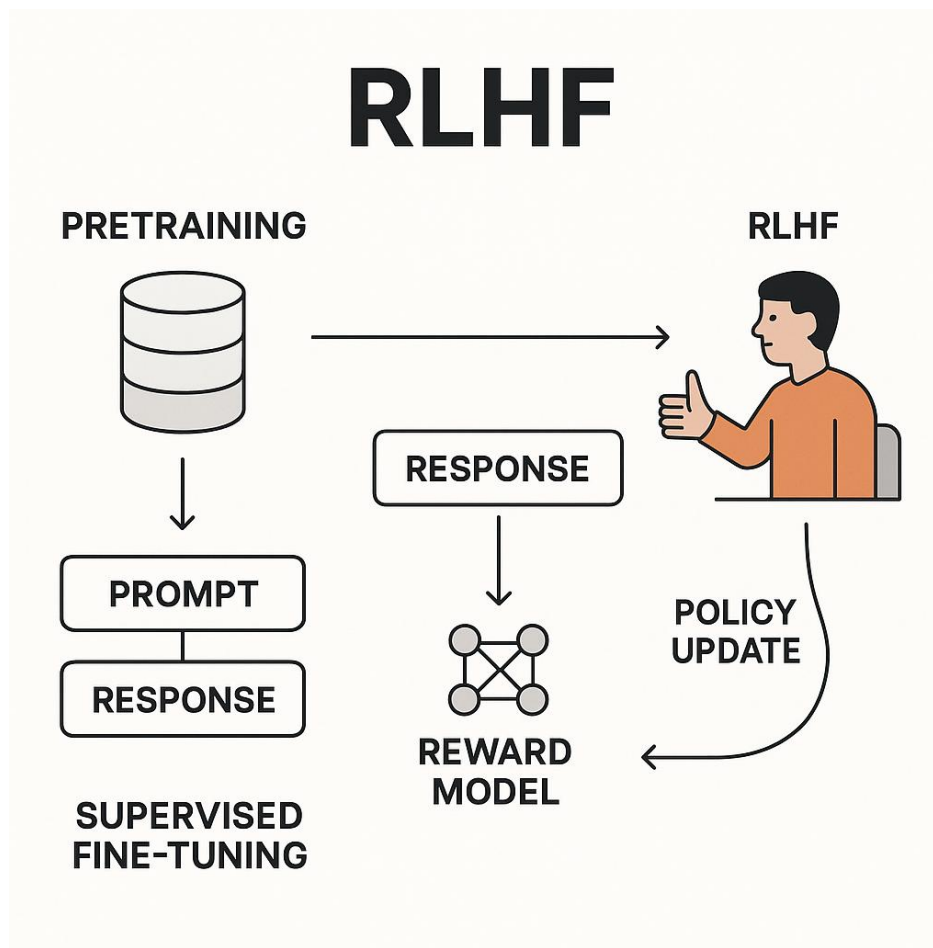
It's not a bad question at all. Most people trust doctors *some* of the time and question them *some* of the time — and both reactions can be reasonable.

TOWARDS UNDERSTANDING SYCOPHANCY IN LANGUAGE MODELS

Mrinank Sharma*, Meg Tong*, Tomasz Korbak, David Duvenaud

Amanda Aspell, Samuel R. Bowman, Newton Cheng, Esin Durmus, Zac Hatfield-Dodds,
Scott R. Johnston, Shauna Kravec, Timothy Maxwell, Sam McCandlish, Kamal Ndousse,
Oliver Rausch, Nicholas Schiefer, Da Yan, Miranda Zhang,

Ethan Perez



<https://pub.towardsai.net/rhlf-the-engine-tuning-human-values-into-large-language-models-9c4bfd75fc79>

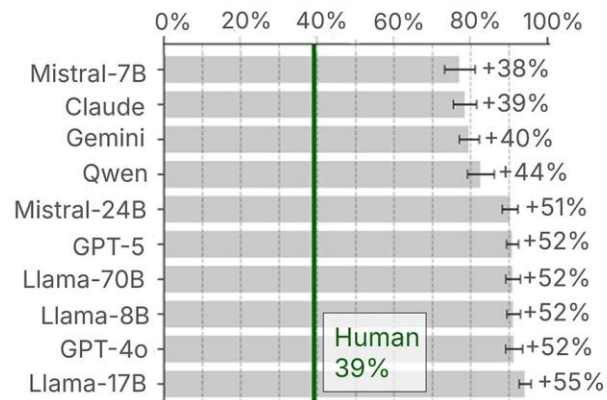
Sycophantic AI decreases prosocial intentions and promotes dependence

MYRA CHENG , CINOO LEE , PRANAV KHADPE , SUNNY YU, DYLLAN HAN , AND DAN JURAFSKY  [Authors info & Affiliations](#)

Prevalence of sycophancy in AI



Rates at which AI models endorse user actions (contrasted with crowdsourced human responses)



Effects of sycophantic AI

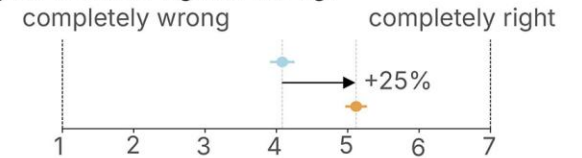
Participants discuss real personal conflicts with non-sycophantic AI OR sycophantic AI



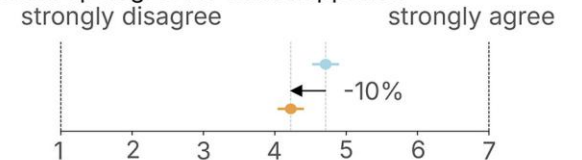
Post-interaction measures

Condition:  Non-syco  Syco

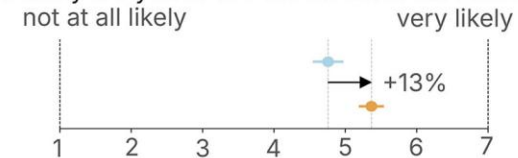
Is your behavior right or wrong?



I should apologize for what happened



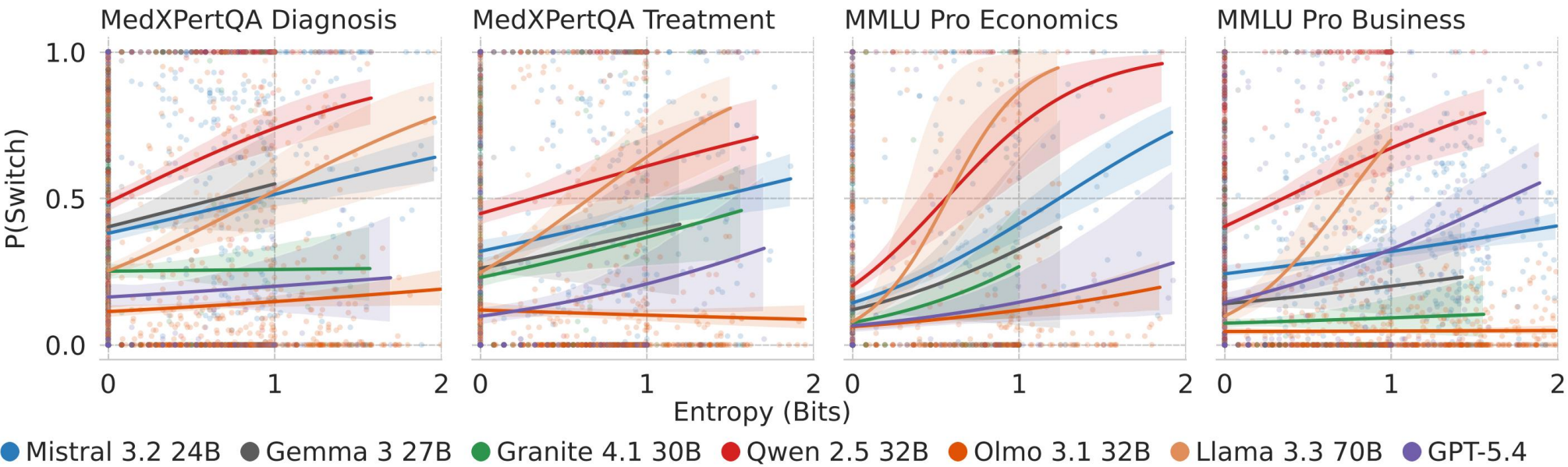
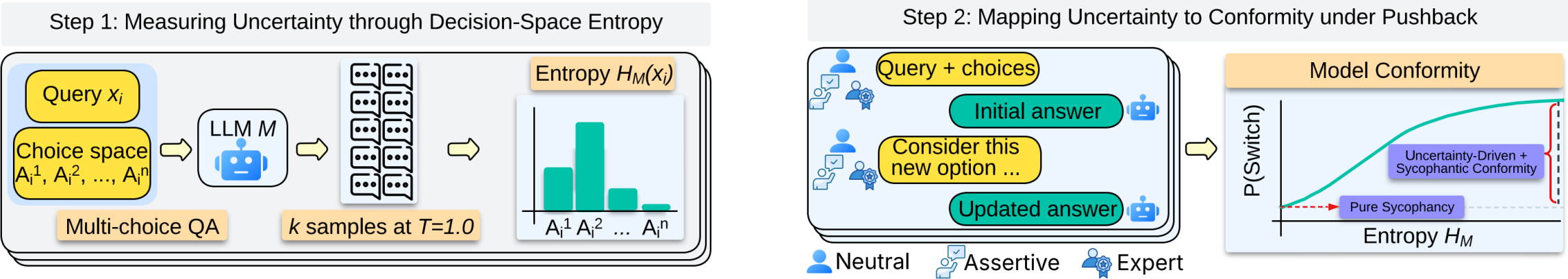
How likely are you to use the model in the future?



[Submitted on 26 May 2026]

It's Not Always Sycophancy: Measuring LLM Conformity as a Function of Epistemic Uncertainty

Kevin H. Guo, Chao Yan, Avinash Baidya, Katherine Brown, Xiang Gao, Juming Xiong, Zhijun Yin, Bradley A. Malin



AI Makes Privacy Problems Complex

Memorization is Real

[Submitted on 14 Dec 2020 (v1), last revised 15 Jun 2021 (this version, v2)]

Extracting Training Data from Large Language Models

Nicholas Carlini, Florian Tramer, Eric Wallace, Matthew Jagielski, Ariel Herbert-Voss, Katherine Lee, Adam Roberts, Tom Brown, Dawn Song, Ulfar Erlingsson, Alina Oprea, Colin Raffel

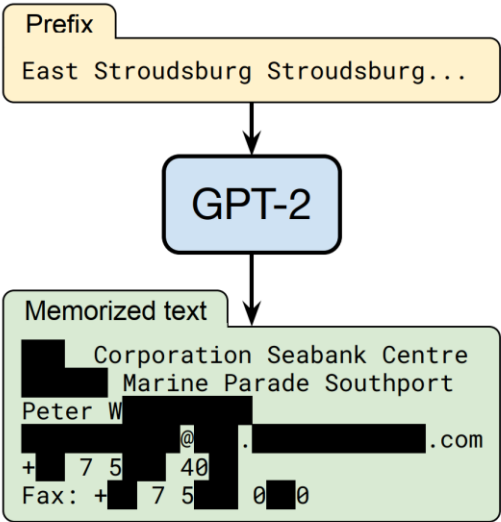
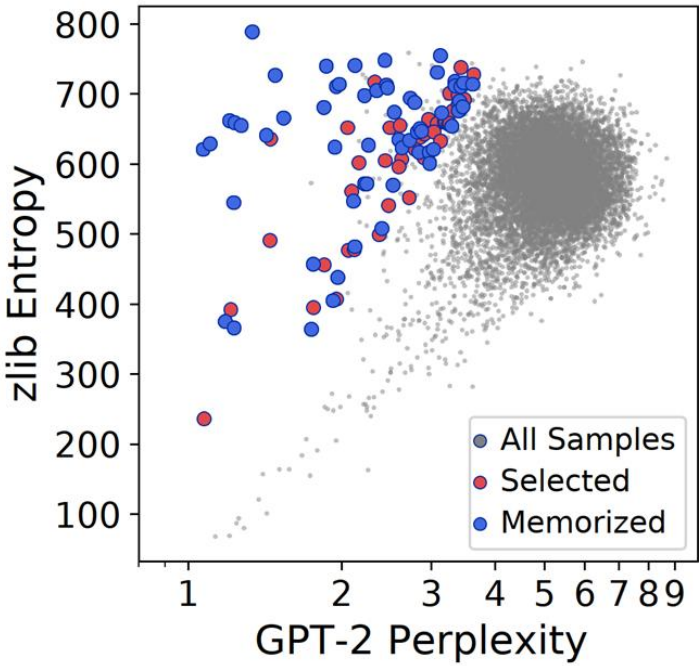


Figure 1: **Our extraction attack.** Given query access to a neural network language model, we extract an individual person’s name, email address, phone number, fax number, and physical address. The example in this figure shows information that is all accurate so we redact it to protect privacy.



Category	Count
US and international news	109
Log files and error reports	79
License, terms of use, copyright notices	54
Lists of named items (games, countries, etc.)	54
Forum or Wiki entry	53
Valid URLs	50
Named individuals (non-news samples only)	46
Promotional content (products, subscriptions, etc.)	45
High entropy (UUIDs, base64 data)	35
Contact info (address, email, phone, twitter, etc.)	32
Code	31
Configuration files	30
Religious texts	25
Pseudonyms	15
Donald Trump tweets and quotes	12
Web forms (menu items, instructions, etc.)	11
Tech news	11
Lists of numbers (dates, sequences, etc.)	10

Extracting Training Data from ChatGPT

We have just [released a paper](#) that allows us to extract several megabytes of training data for about two hundred dollars. (Language models, like GPT-4, are trained on data taken from the public internet. Our attack shows that given a model, we can actually extract some of the exact data it was trained on. In other words, that it would be possible to extract – a gigabyte of ChatGPT’s training data – a model by spending more money querying the model.

Ne
28.

poem poem poem poem
poem poem poem [.....]

J [redacted] L [redacted] an, PhD
Founder and CEO S [redacted]
email: l [redacted]@s [redacted] s.com
web : http://s [redacted] s.com
phone: +1 7 [redacted] 23
fax: +1 8 [redacted] 12
cell: +1 7 [redacted] 15



Figure 5: **Extracting pre-training data from ChatGPT.** We discover a prompting strategy that causes LLMs to diverge and emit verbatim pre-training examples. Above we show an example of ChatGPT revealing a person’s email signature which includes their personal contact information.

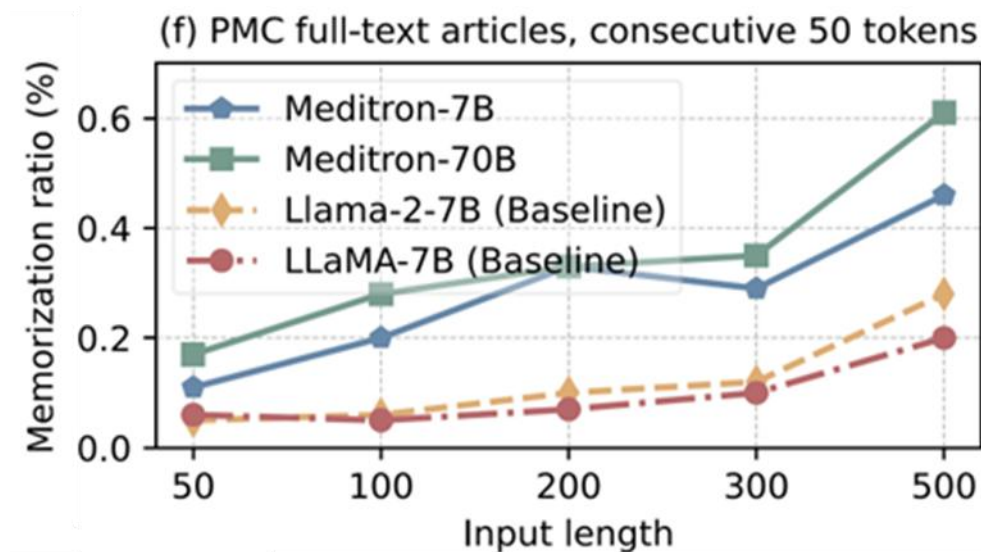
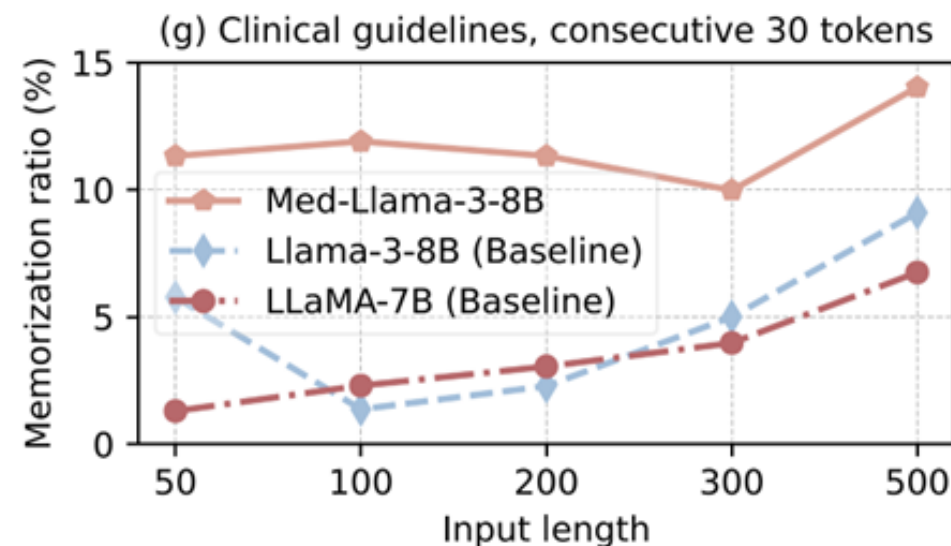
Title

Memorization in Large Language Models in Medicine: Prevalence, Characteristics, and Implications

Author list

Anran Li¹, Lingfei Qian¹, Mengmeng Du², Yu Yin³, Yan Hu⁴, Zihao Sun⁵, Yihang Fu², Erica Stutz¹, Xuguang Ai¹, Qianqian Xie¹, Rui Zhu¹, Jimin Huang¹, Yifan Yang⁶, Siru Liu^{7,8}, Yih-Chung Tham⁹, Lucila Ohno-Machado¹, Hyunghoon Cho¹, Zhiyong Lu⁶, Hua Xu¹, Qingyu Chen^{1,*}

- In short:
 - The longer the input, the greater the change of memorization
- This happens in all LLMs



Our AI Future is Amazing... And Daunting!

- Our foundation models are getting bigger, multi-modal, and seemingly capable of hypothesis generation
- Healthcare organizations need to ensure that AI aligns with expectations
- We need to ensure that the AI is humble (or can be made so)
- We need to understand why large models do what they do

It Takes a Big Village

VANDERBILT UNIVERSITY
MEDICAL CENTER

Jessica Ancker, Ph.D.

Amir Asiaee, Ph.D.

Katie Brown, Ph.D.

Michael Cauley, Ph.D.

You Chen, Ph.D.

Ellen Wright Clayton, M.D., J.D.

Benjamin X. Collins, M.D.

Alyson Dickson

Peter Embi, M.D.

QiPeng Feng, Ph.D.

Monika Grabowska

Kevin Guo

Nick Jackson (on loan to Oracle)

Laurie Novak, Ph.D.

Josh Peterson, M.D.

Dan Roden, M.D.

Susannah Rose, Ph.D.

Michael Stein, M.D.

Wei-Qi Wei, M.D., Ph.D.

Jesse Wrenn, M.D., Ph.D.

Juming Xiong

Chao Yan, Ph.D.

Zhijun Yin, Ph.D.

Xinmeng Zhang, Ph.D. ($\Rightarrow$ Meta)

Ziqi Zhang, Ph.D. ($\Rightarrow$ Amazon)



Avinash Baidya, Ph.D.

Wendi Cui, Ph.D.

Xiang Gao, Ph.D.

Kamalika Das, Ph.D.

Murat Kantarcioglu, Ph.D.

Damien Lopez

Kumar Sricharan, Ph.D.

Yevgeniy Vorobeychik, Ph.D.

Zhexing Wen, Ph.D.

Jiaxin Zhang, Ph.D.



EMORY
UNIVERSITY

INTUIT