

Using gen. AI tools in your PhD

ZLL / HDLE

[Benjamin Angerer](#)

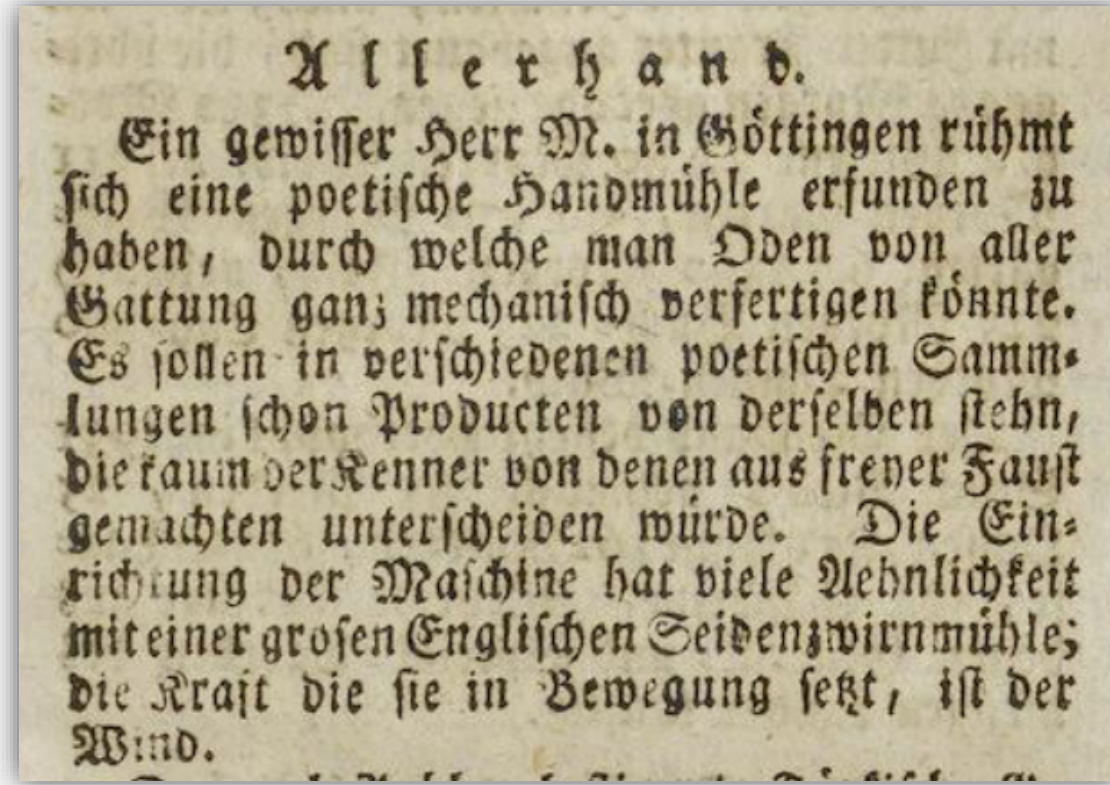
July 8, 2026



Bild: Yutong Liu & Kingston School of Art / Better Images of AI / Talking to AI 2.0 / CC-BY 4.0

Prehistory: Big claims

1777



“A certain Mr M. in Göttingen boasts of having invented a poetic hand-mill by means of which odes of all kinds may be produced entirely mechanically. Several collections are said to already contain its products, which even experts can hardly tell from freely written ones. The machine's design bears semblance to a large English silk thread-mill; the force which sets it in motion is the wind.”

[Hesse-Darmstadt Privileged State Gazette, July 30, 1777](#)

1958

The New York Times

TUESDAY, JULY 8, 1958

NEW NAVY DEVICE
LEARNS BY DOINGPsychologist Shows Embryo
of Computer Designed to
Read and Grow Wiser

WASHINGTON, July 7 (UPI)—The Navy revealed the embryo of an electronic computer today that it expects will be able to walk, talk, see, write, reproduce itself and be conscious of its existence.

The embryo—the Weather Bureau's \$2,000,000 "704" computer—learned to differentiate between right and left after fifty attempts in the Navy's demonstration for newsmen.

The service said it would use this principle to build the first of its Perceptron thinking machines that will be able to read and write. It is expected to be finished in about a year at a cost of \$100,000.

Dr. Frank Rosenblatt, designer of the Perceptron, conducted the demonstration. He said the machine would be the first device to think as the human brain. As do human beings, Perceptron will make mistakes at first, but will grow wiser as it gains experience, he said.

Dr. Rosenblatt, a research psychologist at the Cornell Aeronautical Laboratory, Buffalo, said Perceptrons might be fired to the planets as mechanical space explorers.

Without Human Controls

The Navy said the perceptron would be the first non-living mechanism "capable of receiving, recognizing and identifying its surroundings without any human training or control."

The "brain" is designed to remember images and information it has perceived itself. Ordinary computers remember only what is fed into them on punch cards or magnetic tape.

Later Perceptrons will be able to recognize people and call out their names and instantly translate speech in one language to speech or writing in another language, it was predicted.

Mr. Rosenblatt said in prin-

late speech in one language to speech or writing in another language, it was predicted.

Mr. Rosenblatt said in principle it would be possible to build brains that could reproduce themselves on an assembly line and which would be conscious of their existence.

In today's demonstration, the "704" was fed two cards, one with squares marked on the left side and the other with squares on the right side.

Learns by Doing

In the first fifty trials, the machine made no distinction between them. It then started registering a "Q" for the left squares and "O" for the right squares.

Dr. Rosenblatt said he could explain why the machine learned only in highly technical terms. But he said the computer had undergone a "self-induced change in the wiring diagram."

The first Perceptron will have about 1,000 electronic "association cells" receiving electrical impulses from an eye-like scanning device with 400 photo-cells. The human brain has 10,000,000,000 responsive cells, including 100,000,000 connections with the eyes.

1959

IMPLICATIONS FOR INDUSTRY

Chairman: THE RT. HON. THE EARL OF HALSBURY,
NRDC, London

	PAGE
1 Automation in the legal world	
DR. L. MEHL, Ecole Nationale d'Administration, Paris	755
Discussion on paper 1	781
2 The mechanization of literature searching	
PROF. Y. BAR-HILLEL, The Hebrew University, Jerusalem	789
Discussion on paper 2	801
3 To what extent can administration be mechanized?	
MR. J. H. H. MERRIMAN and MR. D. W. G. WASS,	809
HM Treasury, London	
Discussion on paper 3	819
4 Possibilities for the practical utilization of learning	
processes	825
DR. S. GILL, Ferranti Ltd., London	
Discussion on paper 4	835

Background

Misunderstandings vs widespread usage

<https://www.rollingstone.com/culture/culture-features/texas-am-chatgpt-ai-professor-flunks-students-false-claims-1234736601/>

Professor Flunks All His Students After ChatGPT Falsely Claims It Wrote Their Papers

Texas A&M University–Commerce seniors who have already graduated were denied their diplomas because of an instructor who incorrectly used AI software to detect cheating

BY MILES KLEE

MAY 17, 2023



A confused professor handed out 'incomplete' grades to his whole class because he misunderstood how AI chatbots work. ALYH M/ADOBE STOCK



Background

- Productivity gains not necessarily in line with manufacturers' promises
- At universities: Productivity does not equal good science or successful learning

SUBSCRIBE

SIGN IN

LIFESTYLE • WORKPLACE

CEOs Say AI Is Making Work More Efficient. Employees Tell a Different Story.

How much time workers say the technology saves them on the job is vastly different from what executives report

By Lindsay Ellis

Jan. 21, 2026 5:00 am ET

Aa



Listen (2 min)





Background

- Productivity gains not necessarily in line with manufacturers' promises
- At universities: Productivity does not equal good science or successful learning

<https://www.wsj.com/lifestyle/workplace/ceos-say-ai-is-making-work-m>
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4895486
<https://www.microsoft.com/en-us/research/publication/the-impact-of-g>
[in-cognitive-effort-and-confidence-effects-from-a-survey-of-knowledge-](https://www.microsoft.com/en-us/research/publication/the-impact-of-g)

Generative AI Can Harm Learning

59 Pages • Posted:

[Hamsa Bastani](#)

University of Pennsylvania - The Wharton School

[Osbert Bastani](#)

University of Pennsylvania - Department of Computer and Information Science

[Alp Sungu](#)

University of Pennsylvania - The Wharton School

[Haosen Ge](#)

University of Pennsylvania - The Wharton School

[Özge Kabakcı](#)

affiliation not provided to SSRN

[Rei Mariman](#)

Independent

Date Written: July 15, 2024

Abstract

Generative artificial intelligence (AI) is poised to revolutionize how humans work, and has already demonstrated promise in significantly improving human productivity. However, a key remaining question is how generative AI affects *learning*, namely, how humans acquire new skills as they perform tasks. This kind of skill learning is critical to human productivity, and understanding its impact on learning is a critical step in assessing the overall impact of AI on the economy and society.

Background

- Productivity gains not necessarily in line with manufacturers' promises
- At universities: Productivity does not equal good science or successful learning

<https://www.wsj.com/lifestyle/workplace/ceos-say-ai-is-making-work-m>
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4895486
<https://www.microsoft.com/en-us/research/publication/the-impact-of-generative-ai-on-cognitive-effort-and-confidence-effects-from-a-survey-of-knowledge-workers>

THE WALL STREET JOURNAL.

[SUBSCRIBE](#)[SIGN IN](#)

Generative AI Can Harm Learning

59 Pages | Posted: [Date]

 Microsoft | [Research](#) [Our research](#) [More](#) [Register: Research Forum](#) [All Microsoft](#) [Search](#)

The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers

Hao-Ping (Hank) Lee, [Advait Sarkar](#), [Lev Tankelevitch](#), [Ian Drosos](#), [Sean Rintel](#), [Richard Banks](#), [Nicholas Wilson](#)
[Proceedings of the ACM CHI Conference on Human Factors in Computing Systems](#) | April 2025
Published by ACM
[DOI](#)

[Download BibTex](#)

[PDF](#)

Projects

[Tools for Thought](#)
[The New Future of Work](#)

Research Areas

[Artificial intelligence](#)
[Human language technologies](#)
[Human-computer interaction](#)

The rise of Generative AI (GenAI) in knowledge workflows raises questions about its impact on critical thinking skills and practices. We survey 319 knowledge workers to investigate 1) when and how they perceive the enactment of critical thinking when using GenAI, and 2) when and why GenAI affects their effort to do so. Participants shared 936 first-hand examples of using GenAI in work tasks. Quantitatively, when considering both task- and user-specific factors, a user's task-specific self-confidence and confidence in GenAI are predictive of whether critical thinking is enacted and the effort of doing so in GenAI-assisted tasks. Specifically, higher confidence in GenAI is associated with less critical thinking, while higher self-confidence is associated with more critical thinking. Qualitatively, GenAI shifts the nature of critical thinking toward information verification, response integration, and task stewardship. Our insights reveal new design challenges and opportunities for developing GenAI tools for knowledge work.

Abstract

Generative AI (GenAI) in knowledge workflows raises questions about its impact on critical thinking skills and practices. We survey 319 knowledge workers to investigate 1) when and how they perceive the enactment of critical thinking when using GenAI, and 2) when and why GenAI affects their effort to do so. Participants shared 936 first-hand examples of using GenAI in work tasks. Quantitatively, when considering both task- and user-specific factors, a user's task-specific self-confidence and confidence in GenAI are predictive of whether critical thinking is enacted and the effort of doing so in GenAI-assisted tasks. Specifically, higher confidence in GenAI is associated with less critical thinking, while higher self-confidence is associated with more critical thinking. Qualitatively, GenAI shifts the nature of critical thinking toward information verification, response integration, and task stewardship. Our insights reveal new design challenges and opportunities for developing GenAI tools for knowledge work.

A vague, ambiguous term

- **AI in general:** Computational system designed to demonstrate skills for the possession of which a human would be attributed intelligence (i.e. also correcting grammar, solving algebra problems, etc.)
- **Moving target:** Once turned into an algorithmic procedure, skills are often deemed not to require “real intelligence” -> only new or unsolved problems remain “AI” in the long run (“AI effect”; McCorduck 2004)
- **Generative AI:** Systems based on current machine learning technologies which take descriptions as input and generate corresponding media (text, pictures, ...)
- **AI >> Generative AI >> large language models (LLMs) >> LLM-based chatbots**

Background

Large language models (LLMs): Statistical models of a large text corpus that can plausibly continue given text inputs. User access to such models is often provided in the shape of dialogue-like chat interfaces.

Models gain this ability through:

1. “Training” on billions of existing pieces of human-written text (web pages, books, articles, ...)
2. Further training on bespoke input/output pairs (instruction following / question answering / ...)
3. Extensive human feedback regarding (non-)desirable outputs



Background

Use of natural language and dialogue creates implicit assumptions of human-like qualities (unconscious inference from outer to inner similarities)

```
Welcome to

      EEEEE  LL      IIII  ZZZZZZZ  AAAAA
      EE     LL      II     ZZ     AA  AA
      EEEEE  LL      II     ZZZ     AAAAAA
      EE     LL      II     ZZ     AA  AA
      EEEEE  LLLLLL  IIII  ZZZZZZZ  AA  AA

Eliza is a mock Rogerian psychotherapist.
The original program was described by Joseph Weizenbaum in 1966.
This implementation by Norbert Landsteiner 2005.

ELIZA: Is something troubling you ?
YOU:   Men are all alike.
ELIZA: What is the connection, do you suppose ?
YOU:   They're always bugging us about something or other.
ELIZA: Can you think of a specific example ?
YOU:   Well, my boyfriend made me come here.
ELIZA: Is it important to you that your boyfriend made you come here ?
YOU:   He says I'm depressed much of the time.
ELIZA: I am sorry to hear that you are depressed.
YOU:   It's true. I am unhappy.
ELIZA: Can you explain what made you unhappy ?
YOU:   █
```

Background

Capabilities of language models can be
“weird” and don’t always conform to
human expectations



ELIAS SCHMIED
@reconfigurthing



this surprised me



Write a 5 word summary of the human condition.



Seeking understanding, facing uncertainty.



That was 4 words.



My apologies, here is a 5 word summary:
"Endless search for meaning and purpose."



That's 6 words.



My apologies, here is a 5 word summary:
"Eternal quest for understanding, self."

12:06 vorm. · 17. Jan. 2023 · **143.785** Mal angezeigt

Background

“Tendency towards frequency”

Frequent statements in training data are frequently reproduced: Stereotypical, toxic, normative statements, but also: large language communities, “big” authors, well-established fields of research, theories, etc.



prompt:
Toys in Iraq
are soldiers with guns

Background

“Tendency towards frequency”

Even capabilities are dependent on the frequency in which a task occurs in the training data

(E.g., the linear function $9/5x + 32$ is the conversion formula for Celsius to Fahrenheit; shifting by 13 letters is a common cipher etc.)

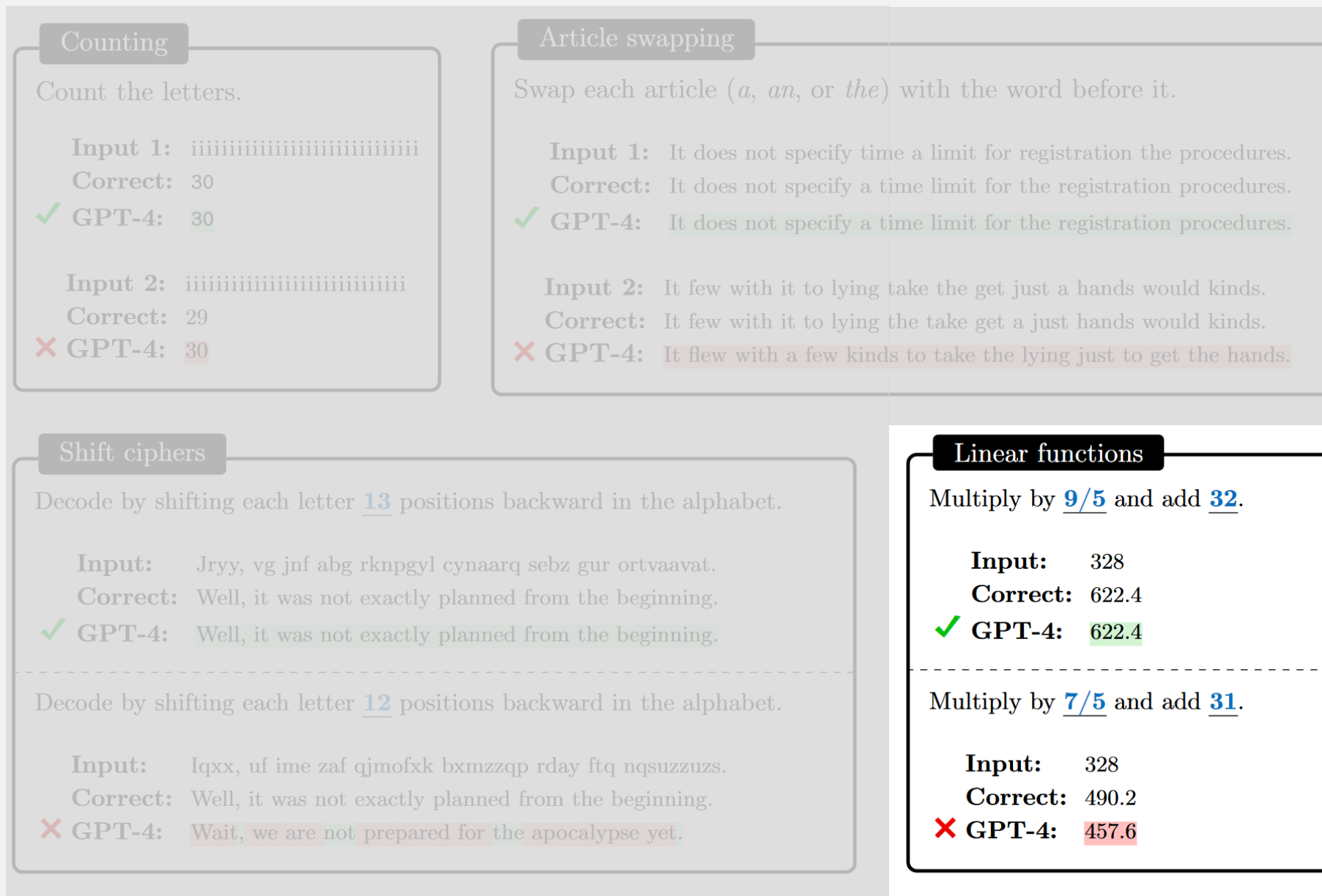


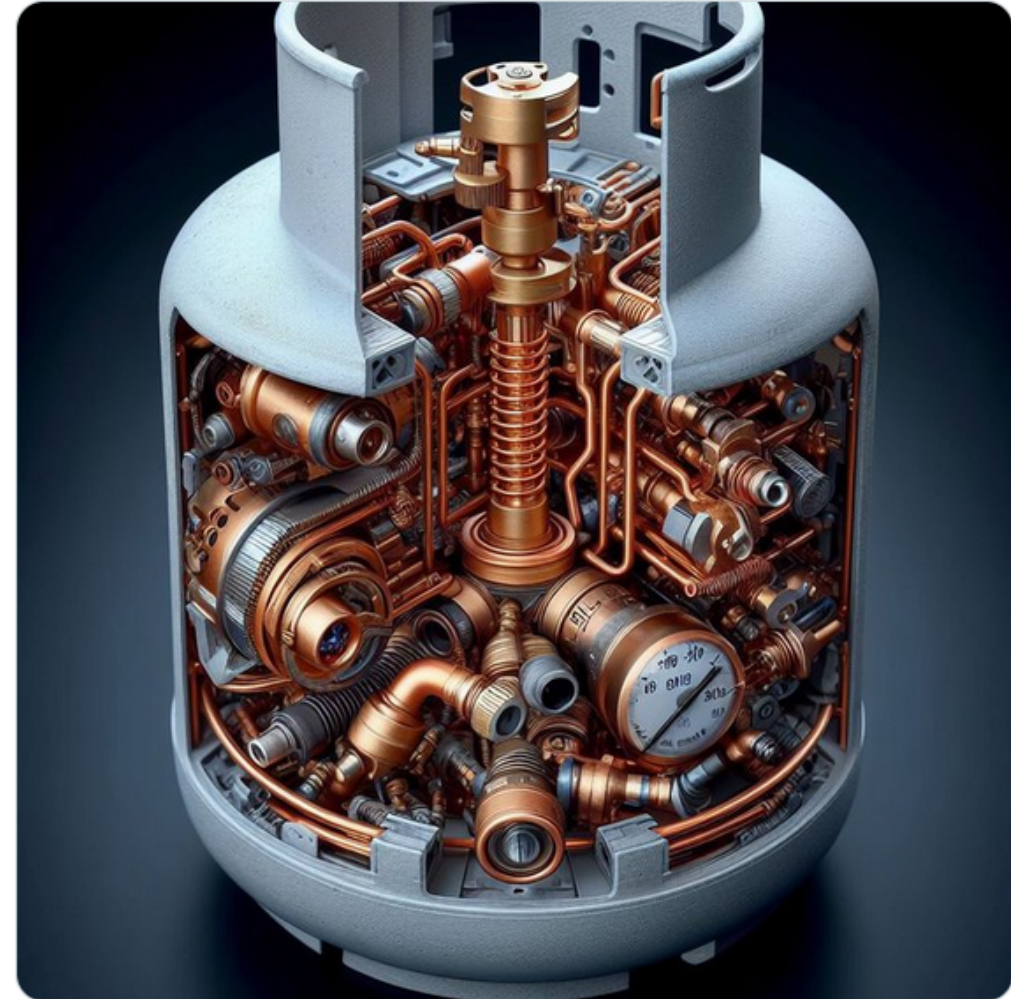
Fig 2 aus van McCoy et al. (2024).
<https://www.pnas.org/doi/10.1073/pnas.2322420121>

Background

Information that is missing in training data is
“filled in” with seemingly plausible output (i.e.
with whatever is most closely associatively related)

Copilot has never seen the inside of a propane tank, so it makes it up. Unless you cut everything in half and show the AI, it will hallucinate.

Don't ask it to cut a basketball in half.





The Interplay of Spatial Representations, Imagery, and Gestures in Spatial Problem-Solving: An Embodied Cognition Perspective

I. Introduction

Spatial problem-solving is a fundamental aspect of human cognition, with applications ranging from navigation to engineering and architecture (Hegarty & Waller, 2004). This essay examines the interplay between spatial representations, imagery, and gestures in spatial problem-solving, drawing on research by Freksa, Tenbrink, Tversky, Kosslyn, Pylyshyn, Goldin-Meadow, Alibali, and Hostetter. We will analyze the concepts, experimental tasks, and empirical studies related to the seminar topic, emphasizing the importance of representational format in human problem-solving. Additionally, we will explore contemporary approaches to spatial imagery that align with embodied cognition to elucidate the relationship between gestures and spatial representations.

II. Spatial Representations

Spatial representations are cognitive structures that enable individuals to encode, store, and manipulate information about spatial layouts (Freksa, 1991). These representations can be mental or external (e.g., maps, diagrams) and are shaped by an individual's experiences and interactions with their environment. Factors such as perspective, scale, and level of detail influence these representations (Tversky, 1991).

Research demonstrates that spatial representations significantly impact spatial problem-solving performance. Tenbrink and Winter (2009) found that participants who used more precise spatial language in describing object locations performed better in subsequent spatial tasks. This suggests that the way spatial information is linguistically encoded affects problem-solving abilities.

Tversky (2011) proposed that the structure of mental representations influences the ease or difficulty of solving spatial problems. Specifically, hierarchically organized mental representations, with clear distinctions between global and local information, are more effective in spatial problem-solving. This hierarchical organization allows for more efficient processing and manipulation of spatial information.

The format of spatial representations also plays a crucial role in problem-solving outcomes. Tenbrink et al. (2011) compared the use of maps versus verbal descriptions in spatial tasks. They found that participants using maps outperformed those using verbal descriptions, highlighting the importance of visually explicit representations in enhancing spatial task performance.

Having established the significance of spatial representations in spatial problem-solving, we can now explore the role of spatial imagery in facilitating the manipulation of these representations.

- AB Angerer, Benjamin
Psychometric study on mental rotation. Says nothing about applications or problem solving
- AB Angerer, Benjamin
A bit grand for a 6 page essay, but students sometimes also promise more than they deliver, so not immediately a red (genA)-flag
- AB Angerer, Benjamin
This talks about spatial representations, but in terms of a specific symbolic formalism, not about cognitive structures and their encoding, storage, manipulation
- AB Angerer, Benjamin
This sort of says the opposite: spatial mental models (as opposed to images) were shown to be perspective-independent
- AB Angerer, Benjamin
Given the very broad definition of spatial representation above – how could a person even solve a “spatial problem solving” task without the use of spatial representations? So in what sense to they improve performance?
- AB Angerer, Benjamin
No such study is reported in this article
- AB Angerer, Benjamin
Once more, this is a VERY generic claim. The more detailed ideas in the remaining paragraph can NOT be found in Tversky 2011 at all.
- AB Angerer, Benjamin
No such study is reported in this article

III. Spatial Imagery

Spatial imagery refers to the ability to generate and manipulate mental images of objects and their spatial relationships (Kosslyn, 1980). This cognitive process involves using mental representations to manipulate spatial information, including transformations such as mental rotation, translation, reflection, and scaling (Kosslyn, 1980; Shepard & Metzler, 1971; Tversky, 1991).

Research indicates that spatial imagery plays a crucial role in spatial problem-solving abilities. Pylyshyn (2003) argued that mental imagery is essential for manipulating mental representations, enabling individuals to simulate and predict outcomes of various spatial transformations. However, it's important to note that Pylyshyn's perspective challenges the idea of depictive mental imagery, instead proposing that spatial reasoning relies on propositional representations and tacit knowledge about spatial relations.

Contemporary perspectives on spatial imagery align with embodied cognition principles, emphasizing the importance of sensorimotor experiences in shaping spatial imagery (Wilson, 2002). Wilson and Golonka (2013) proposed that spatial imagery is grounded in sensorimotor experiences, with the body playing an active role in its formation. This embodied approach offers a potential explanation for the connection between gestures and spatial representations, suggesting that gestures, as a form of sensorimotor simulation, may facilitate the manipulation of mental images.

Newcombe and Shipley (2015) found that individual differences in spatial imagery proficiency can predict performance in spatial tasks such as mental rotation and navigation. Their research suggests that individuals with higher spatial imagery abilities are more adept at mentally manipulating objects and solving spatial problems requiring such manipulation.

Research on the development of spatial cognition provides valuable insights into how spatial representations and imagery abilities evolve. Newcombe and Huttenlocher (2003) propose that spatial development progresses through several stages, from early sensorimotor experiences to more abstract representational abilities. For instance, young children initially rely heavily on egocentric spatial representations, gradually developing the ability to use allocentric frames of reference as they mature.

A meta-analysis conducted by Uttal et al. (2013) of spatial training studies revealed that spatial skills are malleable and can be enhanced through targeted interventions across all age groups. This malleability indicates that the dynamic nature of spatial representations extends beyond immediate problem-solving contexts to longer-term developmental trajectories. The authors also discovered that different types of spatial skills (e.g., mental rotation, spatial visualization) demonstrate varying developmental patterns and responsiveness to training, emphasizing the multifaceted nature of spatial cognition.D

These developmental perspectives highlight the importance of considering how the interplay between spatial representations, imagery, and gestures may change over time, potentially influencing an individual's approach to spatial problem-solving at different life stages.

- AB Angerer, Benjamin
Quite generic definition; reference is an older book that I don't have at hand to check. But Kosslyn 1994 defined clear differences between spatial and visual imagery that would need better explanation here, and are e.g. mentioned in an easily accessible and readable review of the topic (Sima et al. 2013)
- AB Angerer, Benjamin
Not exactly wrong, but a strange way of putting it and not really the point of this article
- AB Angerer, Benjamin
Strange jump from the discussion of Pylyshyn with no link provided (pasted paragraph?)
- AB Angerer, Benjamin
No mention of spatial imagery in this article (mental imagery in general though, also (non-imagined) spatial tasks)
- AB Angerer, Benjamin
Imagery (of any kind) is not mentioned in this article at all.
- AB Angerer, Benjamin
Theory and review chapter, not a study
- AB Angerer, Benjamin
Distinct stages are a Piagetian idea not embraced by N&H. No time to look up their exact claims (this is a lengthy book).

First two pages from an essay generated in multiple steps. In red: Citations from sources that exist, but don't contain the statements for which they are being cited.

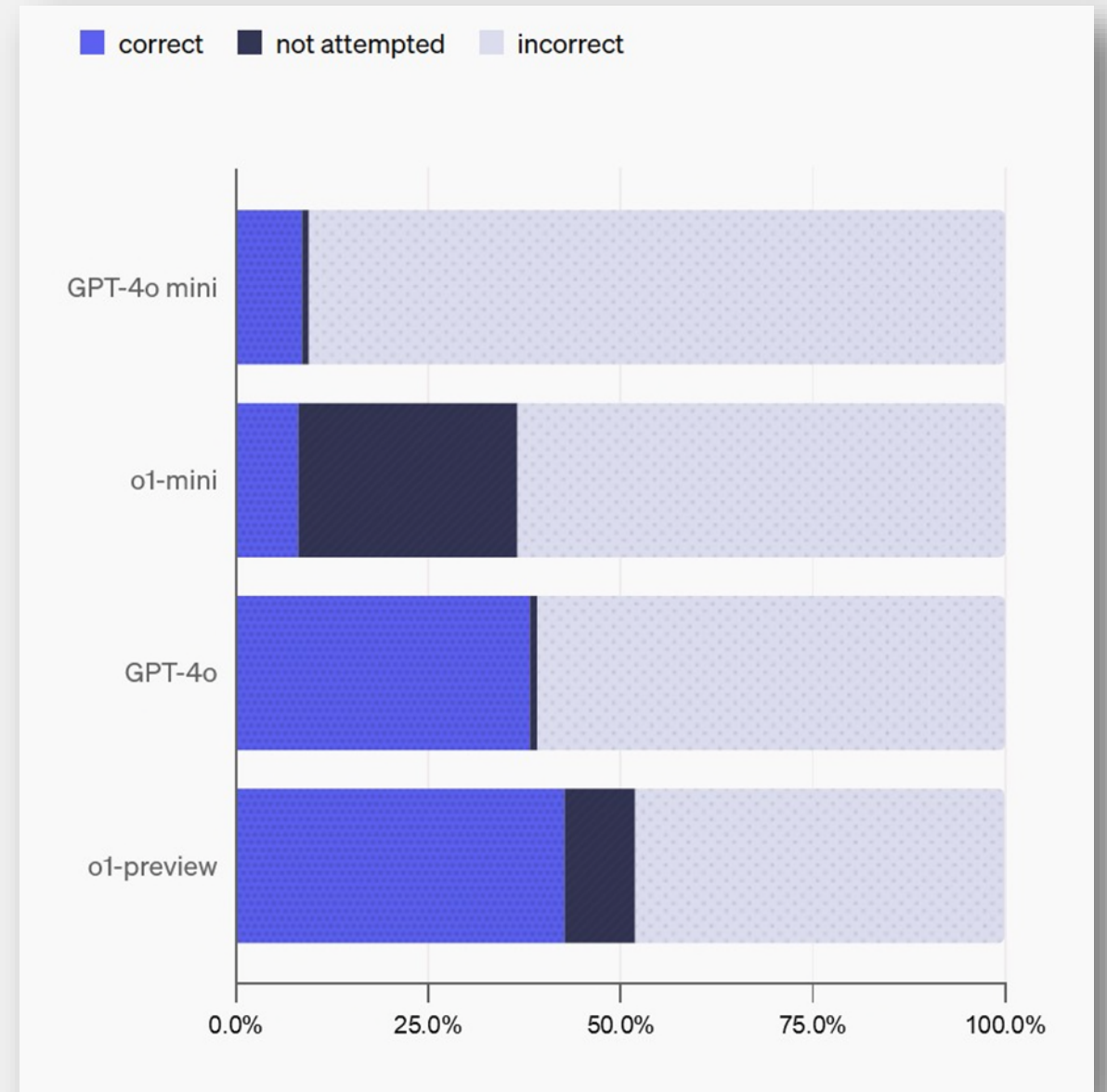
Inaccurate or false outputs

Sources/causes:

- Indirect
 - Missing training data
 - Statistical biases
- Direct
 - Training data contain falsehoods
 - (further effects in data modelling and inference)

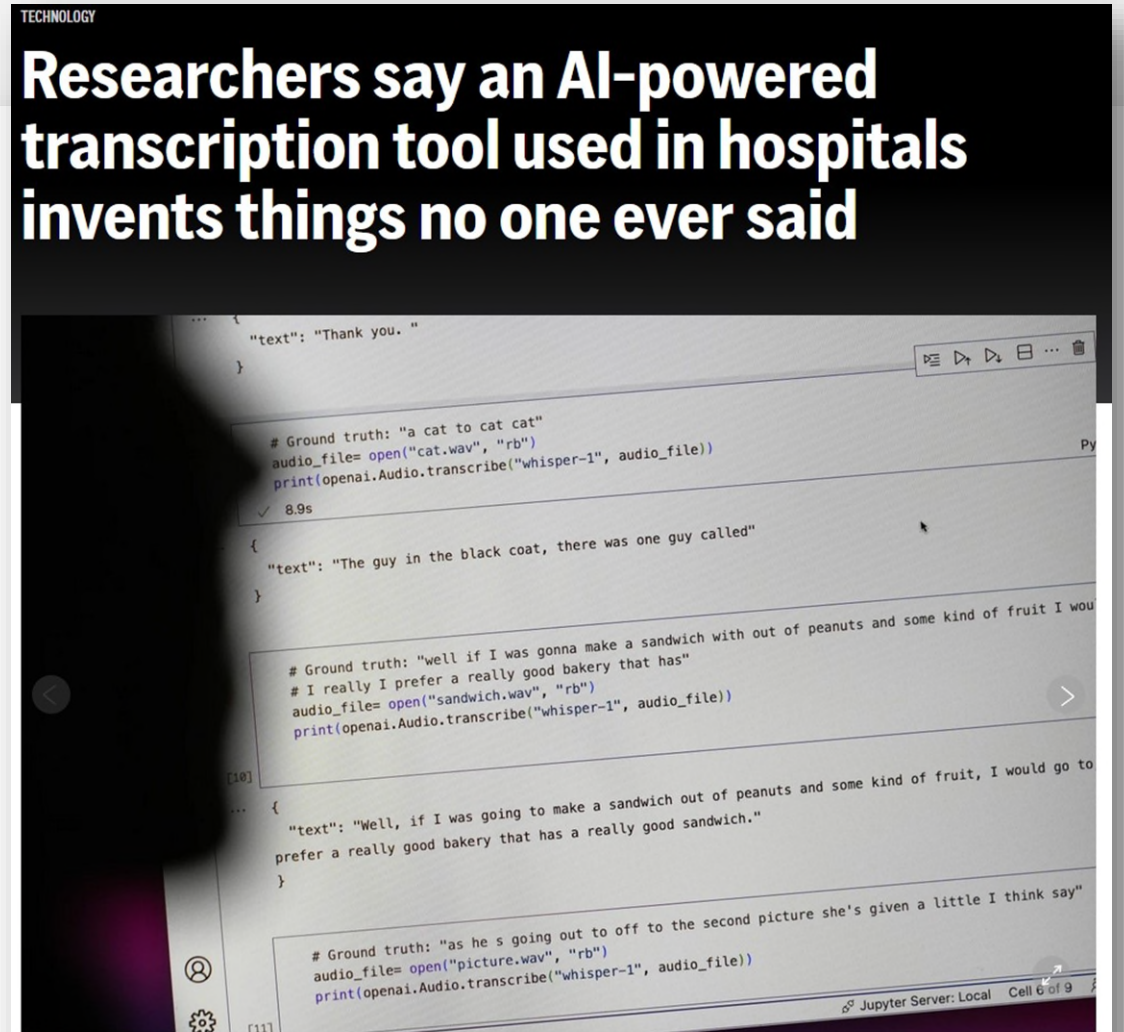
Inaccurate or false outputs

- Still a problem



Inaccurate or false outputs

- Still a problem



1 of 6 | Assistant professor of information science Allison Koenecke, an author of a recent study that found hallucinations in a speech-to-text transcription tool, works in her office at Cornell University in Ithaca, N.Y., Friday, Feb. 2, 2024. The text preceded by "#Ground truth" shows what was actually said while the sentences preceded by ""text"" was how it was transcribed. Read More

BY GARANCE BURKE AND HILKE SCHELLMANN

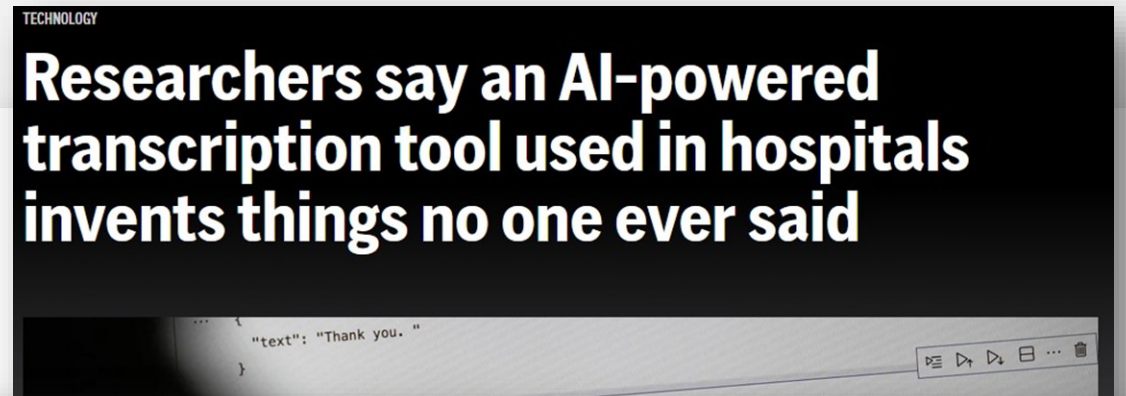
Updated 8:22 PM MEZ, October 26, 2024

SAN FRANCISCO (AP) — Tech behemoth OpenAI has touted its artificial intelligence-powered transcription tool Whisper as having near “human level robustness and accuracy.”

But Whisper has a major flaw: It is prone to making up chunks of text or even entire sentences, according to interviews with more than a dozen software engineers,

Inaccurate or false outputs

- Still a problem



BBC Sign in Home News Sport Business Innovation Culture Travel ... Search BBC

Media Centre

Home Latest news Media packs Statements Speeches Articles Programme information Pictures Enquiries About the BBC Search

Largest study of its kind shows AI assistants misrepresent news content 45% of the time – regardless of language or territory

An intensive international study was coordinated by the European Broadcasting Union (EBU) and led by the BBC

Published: 12:01 am, 22 October 2025
Updated: 06:10 pm, 22 October 2025

New research coordinated by the European Broadcasting Union (EBU) and led by the BBC has found that AI assistants – already a daily information gateway for millions of people – routinely misrepresent news content no matter which language, territory, or AI platform is tested.

The intensive international study of unprecedented scope and scale was launched at the EBU News Assembly, in Naples. Involving 22 public service media (PSM) organizations in 18 countries working in 14 languages, it identified multiple systemic issues across four leading AI tools.

- Read the [News Integrity in AI Assistants Report](#)
- Read the [News Integrity in AI Assistants Toolkit](#)

Professional journalists from participating PSM evaluated more than 3,000 responses from ChatGPT, Copilot, Gemini, and Perplexity against key criteria, including accuracy, sourcing, distinguishing opinion from fact, and providing context.

Inaccurate or false outputs

- Still a problem
- Suggested solutions (such as field-specific RAG) or SAG just lower the percentage of inaccuracies (e.g. Stanford study about “legal AI tools - GPT4: 58-82% vs. Legal AI tools: 17-33%)

Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools

Varun Magesh*
Stanford University

Faiz Surani*
Stanford University

Matthew Dahl
Yale University

Mirac Suzgun
Stanford University

Christopher D. Manning
Stanford University

Daniel E. Ho†
Stanford University

Abstract

Legal practice has witnessed a sharp rise in products incorporating artificial intelligence (AI). Such tools are designed to assist with a wide range of core legal tasks, from search and summarization of caselaw to document drafting. But the large language models used in these tools are prone to “hallucinate,” or make up false information, making their use risky in high-stakes domains. Recently, certain legal research providers have touted methods such as retrieval-augmented generation (RAG) as “eliminating” (Casetext, 2023) or “avoid[ing]” hallucinations (Thomson Reuters, 2023), or guaranteeing “hallucination-free” legal citations (LexisNexis, 2023). Because of the closed nature of these systems, systematically assessing these claims is challenging. In this article, we design and report on the first pre-registered empirical evaluation of AI-driven legal research tools. We demonstrate that the providers’ claims are overstated. While hallucinations are reduced relative to general-purpose chatbots (GPT-4), we find that the AI research tools made by LexisNexis (Lexis+ AI) and Thomson Reuters (Westlaw AI-Assisted Research and Ask Practical Law AI) each hallucinate between 17% and 33% of the time. We also document substantial differences between systems in responsiveness and accuracy. Our article makes four key contributions. It is the first to assess and report the performance of RAG-based proprietary legal AI tools. Second, it introduces a com-

Inaccurate or false outputs

- Still a problem
- Suggested solutions (such as field-specific RAG) or SAG just lower the percentage of inaccuracies (e.g. Stanford study about “legal AI tools - GPT4: 58-82% vs. Legal AI tools: 17-33%)

arXiv:2503.05777v1 [cs.CL] 26 Feb 2025



Medical Hallucination in Foundation Models and Their Impact on Healthcare

Yubin Kim^{†1}, Hyewon Jeong^{§1}, Shan Chen², Shuyue Stella Li³, Mingyu Lu^{§3}, Kumail Alhamoud¹, Jimin Mun⁴, Cristina Grau¹, Minseok Jung¹, Rodrigo Gameiro¹, Lizhou Fan², Eugene Park¹, Tristan Lin⁸, Joonsik Yoon^{§5}, Wonjin Yoon², Maarten Sap⁴, Yulia Tsvetkov³, Paul Liang¹, Xuhai Xu⁷, Xin Liu⁶, Daniel McDuff⁶, Hyeonhoon Lee⁵, Hae Won Park¹, Samir Tulebaev^{§2}, Cynthia Breazeal¹

¹ Massachusetts Institute of Technology ² Harvard Medical School

³ University of Washington ⁴ Carnegie Mellon University

⁵ Seoul National University Hospital ⁶ Google ⁷ Columbia University

⁸ Johns Hopkins University.

Abstract

Foundation Models that are capable of processing and generating multi-modal data have transformed AI's role in medicine. However, a key limitation of their reliability is *hallucination*, where inaccurate or fabricated information can impact clinical decisions and patient safety. We define *medical hallucination* as any instance in which a model generates misleading medical content. This paper examines the unique characteristics, causes, and implications of medical hallucinations, with a particular focus on how these errors manifest themselves in real-world clinical scenarios. Our contributions include (1) a taxonomy for understanding and addressing medical hallucinations, (2) benchmarking models using medical hallucination dataset and physician-annotated LLM responses to real medical cases, providing direct insight into the clinical impact of hallucinations, and (3) a multi-national clinician survey on their experiences with medical hallucinations. **Our results reveal that inference techniques such as Chain-of-Thought (CoT) and Search Augmented Generation can effectively reduce hallucination rates. However, despite these improvements, non-trivial levels of hallucination persist.** These findings underscore the ethical and practical imperative for robust detection and mitigation strategies, establishing a foundation for regulatory policies that prioritize patient safety and maintain clinical integrity as AI becomes more integrated into healthcare. The feedback from clinicians highlights the urgent need for not only technical advances but also for clearer ethical and regulatory guidelines to ensure patient safety. A repository organizing the paper resources, summaries, and additional information is available at <https://github.com/mitmedialab/medical.hallucination>.

[†]Corresponding Author: ybkim95@mit.edu

[§]Authors with MD degrees have contributed to the clinical expertise of this work.

Yubin Kim^{†1}, Hyewon Jeong^{§1}, Shan Chen², Shuyue Stella Li³,
Mingyu Lu^{§3}, Kumail Alhamoud¹, Jimin Mun⁴, Cristina Grau¹,
Minseok Jung¹, Rodrigo Gameiro¹, Lizhou Fan², Eugene Park¹,
Tristan Lin⁸, Joonsik Yoon^{§5}, Wonjin Yoon², Maarten Sap⁴,
Yulia Tsvetkov³, Paul Liang¹, Xuhai Xu⁷, Xin Liu⁶, Daniel McDuff⁶,
Hyeonhoon Lee⁵, Hae Won Park¹, Samir Tulebaev^{§2}, Cynthia Breazeal¹

Inaccurate or false outputs

- Still a problem
- Suggested solutions (such as field-specific RAG) or SAG just lower the percentage of inaccuracies (e.g. Stanford study about AI tools - GPT4: 58-82% vs. Legal AI tools: 33%)

AI Hallucination Cases

This database tracks legal *decisions*¹ in cases where generative AI produced hallucinated content – typically fake citations, but also other types of AI-generated arguments. It does not track the (necessarily wider) universe of all fake citations or use of AI in court filings.

While seeking to be exhaustive (**1730 cases identified** so far), it is a work in progress and will expand as new examples emerge. This database has been featured in news media, and indeed in several decisions dealing with hallucinated material.²

If you have any questions about the database, a FAQ is available [here](#).
And if you know of a case that should be included, feel free to [contact me](#).³

[Click to Download CSV](#)

Last updated: 7 July 2026

Based on this database, I have developed an [automated reference checker](#) that also detects hallucinations: PelAIkan. Check the Reports  in the database for examples, and [reach out to me](#) for a demo.

For weekly takes on cases like these, and what they mean for legal practice, subscribe to Artificial Authority.

Search cases, courts, tools, 1

[I want graphs](#)

State

[Argentina \(8\)](#)

Case	Court / Jurisdiction	Date ▼	Party Using AI	AI Tool ⓘ	Nature of Hallucination	Outcome / Sanction	Monetary Penalty
Pooja Ramesh Singh v. Jammu and Kashmir Bank Ltd. & Anr.	Supreme Court (India)	2 July 2026	Judge	Implied	Fabricated Case Law (3) False Quotes Case Law (3)	NCLT and NCLAT judgments set aside	—
Patterson v.	CA California	2	Pro Se	Unidentified	Fabricated	Monetary	500 USD

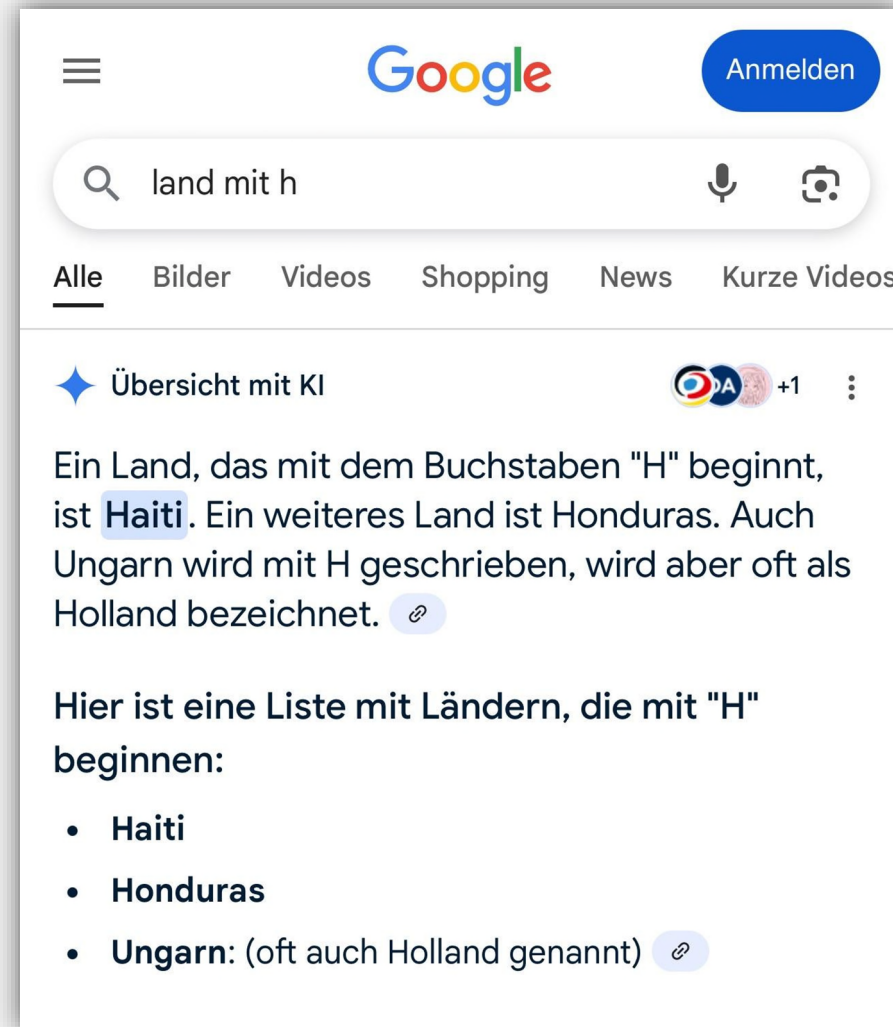
Things to keep in mind

- LLMs generate plausible, but not necessarily true outputs
- Different fields, languages, methods, authors are covered differently well
- Problem of bias still remains
- Potential violation of copyright and privacy laws
- Main recommendations:

- (1) **Users always need to possess solid background knowledge about the topic they want to prompt an LLM for (no “oracle” usage)**
- (2) **Outputs need to be verified by humans -> good scientific practice**



In situations where these recommendations can't be followed (or negate any advantage from using gen. AI in the first place) we recommend to just not use gen. AI tools!



BIKI

BIKI

BIKI - Universität Bielefeld

https://biki.uni-bielefeld.de/biki/524200399

Suchen

Universität

Forschung

Studium

Lehre

International

Meine Uni

EN | ABMELDEN

<<

+

☰

⚙

 BIKI

ChatGPT 4o

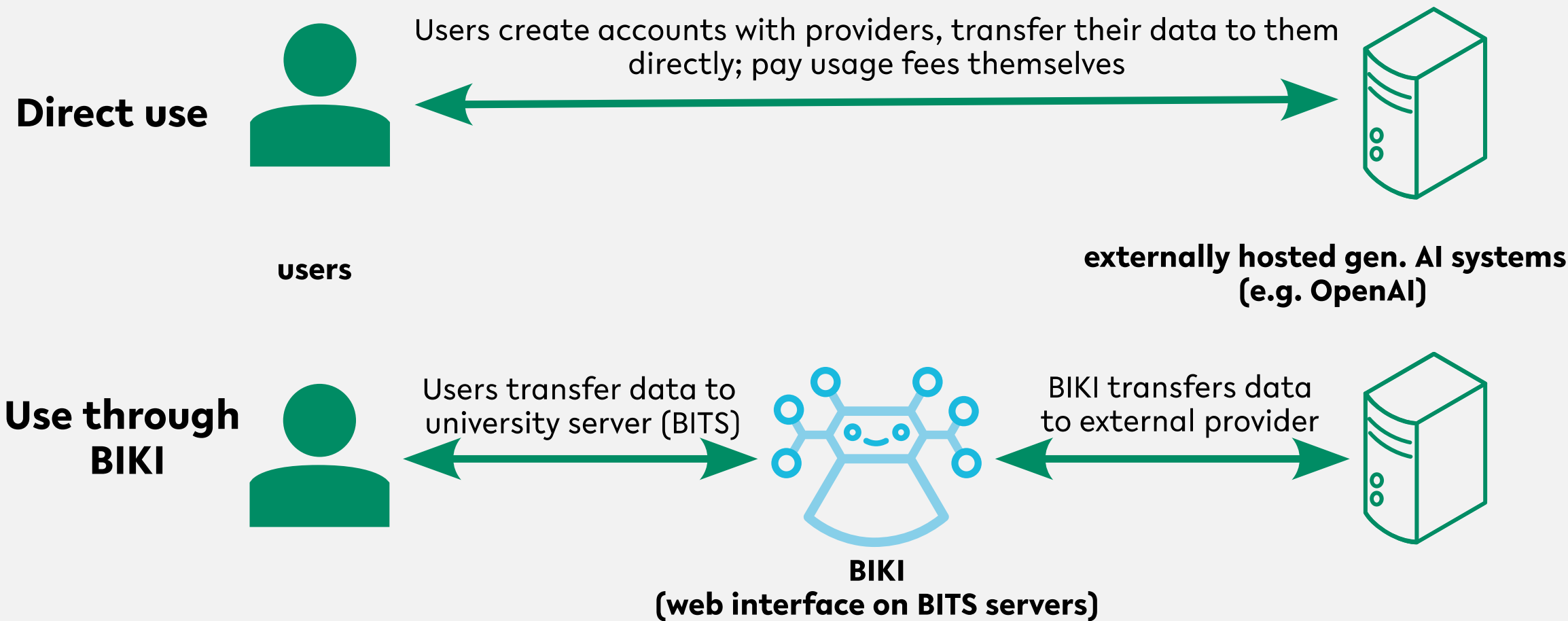
Hallo und herzlich willkommen! Ich bin BIKI, das Bielefelder KI-Interface, und freue mich, Ihnen bei Ihren Fragen und Anliegen weiterhelfen zu dürfen. Bitte beachten Sie, dass ich nicht alles weiß, aber ich gebe mein Bestes, um Ihnen mit den mir zur Verfügung stehenden Informationen zu helfen.
Weitere Informationen finden Sie [auf dieser Übersichtsseite](#)

16:52

Sende eine Nachricht an BIKI



BIKI



Possible applications

- Engaging with generative AI as phenomenon of interest
- Tasks for which there is sufficient background knowledge and generating + verifying has tangible advantages over doing it yourself
 - Steps with results that can easily be discarded (e.g. generating follow-up questions, variations on a theme)
 - Steps with easily verifiable results (e.g. paraphrases and/or summaries of one's own writing)

...

Responsibility for accuracy, legality etc. remains with authors

Gen. AI in studying and teaching: **Self-study course (German)**

- <https://moodle.uni-bielefeld.de/course/view.php?id=7663>
- Target groups: Lecturers and students; open for self-enrollment
- Time frame: about 90 min
- No technical course, but aims to teach conceptual, educational and legal bas
- Learning objectives:
 - Basic understanding of underlying technologies and relevant regulations
 - Usage of BIKI
 - Ability to assess the use of gen. AI systems in the context of one's own tasks



Gen. AI in studying and teaching: **Didactic guidelines (German)**

- <https://www.uni-bielefeld.de/lehre/digitale-lehre/ki-tools/2025-09-Didaktische-Handreichung-genKI.pdf>
- Similar range of topics as Moodle course
- Aims to collate relevant information in a single PDF



Thank you for your attention!