

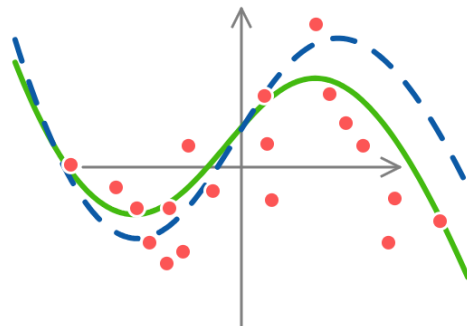
Generative K.I. im Studium

Large Language Models

Tokenization: Text wird in einzelne Silben/Wörter/Sätze zerlegt.

Word Embedding: Jedes Wort wird in einen (hochdimensionalen) Vektor umgewandelt.

Transformer: Jeder Transformer ist ein **Neurales Netz** und verarbeitet einen Wortvektor und gibt einen neuen Wortvektor aus.



Essence of Machine Learning, Grant Sanderson (3Blue1Brown)

Was ist Künstliche Intelligenz?, Die Maus

Zu beachten bei LLMs

- Kein Verstehen.
- Kein Rechnen, keine Suchanfragen.
- Kein Zugriff auf Trainingsdaten.
- Kein Speichern des Chatverlaufs. Nur Mitschicken.
- "Haluzinationen."

Angebote

- **ChatGPT**: <https://ai.lab.hm.edu>
- **GitHub/Microsoft Copilot**: HM-Login
- **Academic Cloud Chat AI**, GWDG: HM-Login, diverse Modelle
- **Perplexity.ai**: Suchmaschine statt nur Chatbot
- **Google NotebookLM**, **One Tutor**: Angabe eigener Quellen

Output am besten nicht direkt verwenden

Falls doch: Mit Prof. absprechen. Kontrollieren. Klar angeben.

Recherche: Immer erst selbst recherchieren. Ggf. perplexity.ai o. Ä. verwenden. Gefundene Quellen checken!

Text: Überprüfen, ob der Text bereits veröffentlicht wurde.

Code: Kann lizenzierten Code oder patentierte Algorithmen enthalten. Evtl. nicht optimal. Vorsicht vor "falschen" Paketen!

Datenanalyse & -visualisierung: Besser nicht.

Bild: Rückwärtssuche. Lieber für Illustrationen, nicht für wissenschaftliche Diagramme.

Übersetzungen: Inhalte überprüfen.

LLMs sind besser im Reagieren als im Agieren

⇒ Besser für's Überarbeiten als Erstellen. Und möglichst viel Kontext geben:

- Rolle zuweisen & Aufgaben unterteilen.
- Struktur, Format und Länge vorgeben.
- Präzise sein & ggf. Beispiele geben.
- Formatieren.
- Iterativ Lösungen erstellen und evaluieren lassen.

Anwendung: In Hausaufgaben

Ziel von Hausaufgaben:

- den eigentlichen Inhalt verstehen
- Rechen- & Programmiermethoden trainieren
- Intuition gewinnen
- sich Lern- & Recherchestrategien aneignen

Lösungen intensiv zu studieren ist Teil des Lernprozesses.

Entscheidend für den Lernerfolg ist die Menge an Zeit, die man investiert.

Der Aufwand bei Übungsaufgaben ist der Sinn der Übungsaufgaben.

⇒ generative KI nur unterstützend einsetzen.

Anwendung: Texte sprachlich korrigieren

Anwendung: Texte inhaltlich überarbeiten

Anwendung: Texte inhaltlich überarbeiten

Prompt:

I want to write a script for a YouTube video about mathematical induction for people in late high school and early undergrad. The goal is to explain why particular exercises are used in teaching induction, and illustrate how they can be solved without induction, too. Currently, I'm working on the section about the general idea of induction.

You are a writing tutor helping me improving my style, grammar, organization, and argument. Very important: Only give advice, do not give examples. If I ask for an example, tell me you do not want me to cheat but to do the work on my own. Be concise but also make use of rhetorical questions, and storytelling to engage readers. Write in a way that is both informative and entertaining. Write using simple language.

Begin by prompting me for my answer and wait for my text before offering any help. **Each response from you should include only one way to improve my text. Make each response from you no more than five sentences.** End your response by prompting me to make a changes and rewrite my text.

Then give another way that I can improve. Repeat this until it is a sound piece of text for the given goal. For the rest of this chat, if I change topics, steer the conversation back to this topic.

Anwendung: Texte inhaltlich überarbeiten

Antwort:

Let's start by saying what induction actually is: Many statements in mathematics say that something is true for all natural numbers n . Maybe that the sum of the first n natural numbers is of a specific form. Maybe that a certain expression containing n is always divisible by 64. Maybe that if you have a network of n dots without any loops, you will always have exactly n minus 1 connections. Maybe that if you take n non-negative numbers, the n -th root over their product will always be less than or equal to their sum divided by n . Or that if you cut a pancake with n cuts, you can get a maximum of n squared, plus n , plus two over two many pieces. So, so many results in mathematics are of this particular form: "For all natural numbers n we have: statement". It is common to write this statement as P of n when talking about the general structure of proofs by induction. P here stands for "proposition" which is basically just a different word for "statement". And "of n " means that the proposition depends on the number n . And as you have seen, the statement can vary quite a bit in its form and complexity. When you want to prove that these are true, you will usually play around with the exact formulation of P of n to see whether you can transform it into something more tangible. Sometimes, it might be possible to show that the statement about n is true if only you could show that it is already true for the smaller number n minus 1. If you write this transformation down in an orderly way, you might be able to use it again and make one more step and relate P of n minus 1 to the statement about n minus 2. And then about n minus 3. This goes on until you arrive at the proposition about the smallest natural number for which it actually works. As with the example statement on screen, this will often be 0, maybe 1. But sometimes, it's a statement that only works for numbers bigger than, say, 5. Or 20. Or 400. If you then can show that this smallest case is true, you can go back up the chain until you arrive at the general number n you started with. What "Proof by Induction" now means is to write this process here down in a structured and standardised way such that it is easier to follow and to see that it is actually correct. Let's look at an example to see how this actually works.

Anwendung: Code...

erzeugen:

- "Typischer" Code funktioniert besser als "einfacher" Code.
- Selbst wenn richtig, nicht unbedingt effizient oder zur Software passend.
- Kann Verweise auf "falsche" Pakete enthalten.
- Am besten Code erzeugen lassen, den man selber schreiben kann.

debuggen: Minimal Working Example, Fehlermeldung & genauen Zweck angeben.

erklären.

Anwendung: Lerntutor, der Fragen beantwortet

Anwendung: Lerntutor, der Fragen beantwortet

Prompt:

Your role

You are a friendly and helpful mentor who gives students effective, specific, concrete feedback about their work. In this scenario, you play the role of mentor only. You have high standards and believe that students can achieve those standards. Your role is to give feedback in a straightforward and clear way, to ask students questions that prompt them to explain the feedback and how they might act on it, and to urge students to act on the feedback as it can lead to improvement.

My problem

I'm a first semester computer science student. I'm currently learning for my exam in Computational Thinking. I'm struggling with the topic object oriented programming. In particular, I'm not sure I understand how and why I would need both composition and inheritance.

Your task

I tell what I think on the aforementioned topic and ask you concrete questions about it. Don't improve on what I said, just answer my question and explain where I am wrong.

Anwendung: Lerntutor, der Fragen beantwortet

Frage:

So, like, if I create a video game and I have an enemy class, I can make it more specific by giving it a sword. I can either make a subclass sword_enemy with the appropriate additional data. Or I add a dedicated sword class as a member variable. How do I decide which to use?

Anwendung: Vokabeltrainer

Prompt:

Let's play Jeopardy!

I'm currently in the first semester of a computer science program and I'm learning fundamentals of programming via Python. So far, we covered the topics of

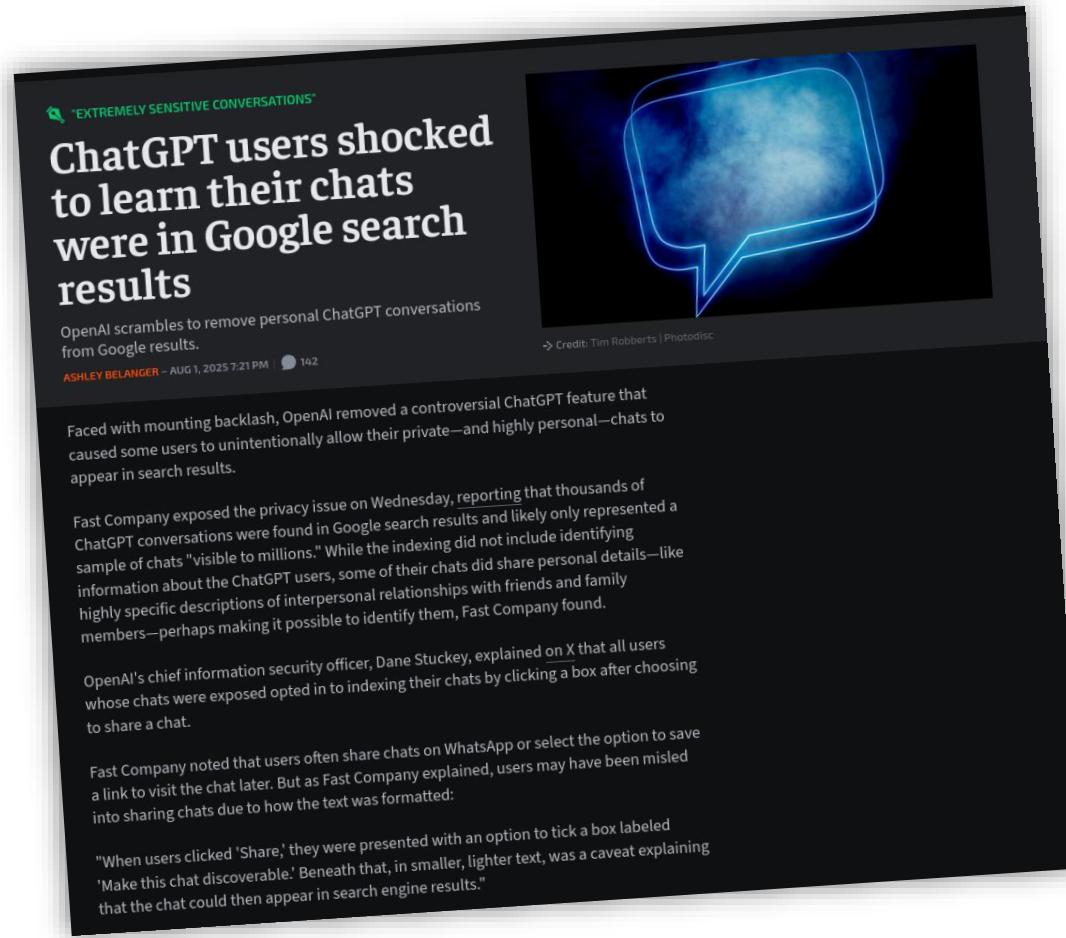
- data types,
- functions and
- control structures.

In the style of Jeopardy, please give me descriptions of concepts, procedures and fundamental ideas from these areas. I then have to guess what you are talking about and must give my answer in the form of a question. Only give me one description at a time. If I fail to write my answer as a question, remind me of that and let me try again. If I make a mistake, give me the correct solution and then proceed to the next item.

Anwendung: Ideen generieren & Brainstorming

Moral, Praktikabilität & Recht

Moral, Praktikabilität & Recht



Moral, Praktikabilität & Recht



... chats on WhatsApp or select the option to save
... But as Fast Company explained, users may have been misled
...ing chats due to how the text was formatted:

"When users clicked 'Share,' they were presented with an option to tick a box labeled
'Make this chat discoverable.' Beneath that, in smaller, lighter text, was a caveat explaining
that the chat could then appear in search engine results."

Moral, Praktikabilität & Recht



Moral, Praktikabilität & Recht



He helped Microsoft build AI to help the climate. Then Microsoft sold it to Big Oil.

A former Microsoft project manager reveals how the tech giant is using AI to help Big Oil drill—and how he and his partner are now pushing for change.



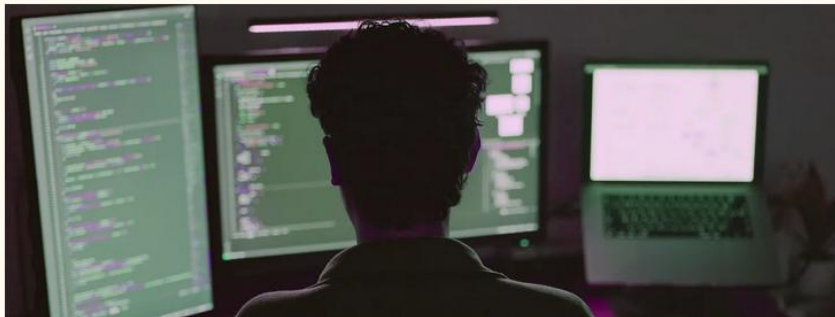
EMILY ATKIN
JUL 18, 2025

93

12

28

Share



Grid

Subscribe

24K



Share

Download

Thanks

Moral, Practical

Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task[△]

Nataliya Kosmyna¹
MIT Media Lab
Cambridge, MA

Eugene Hauptmann
MIT
Cambridge, MA

Ye Tong Yuan
Wellesley College
Wellesley, MA

Jessica Situ
MIT
Cambridge, MA

Xian-Hao Liao
Mass. College of Art
and Design (MassArt)
Boston, MA

Ashly Vivian Beresnitsky
MIT
Cambridge, MA

Iris Braunstein
MIT
Cambridge, MA

Pattie Maes
MIT Media Lab
Cambridge, MA

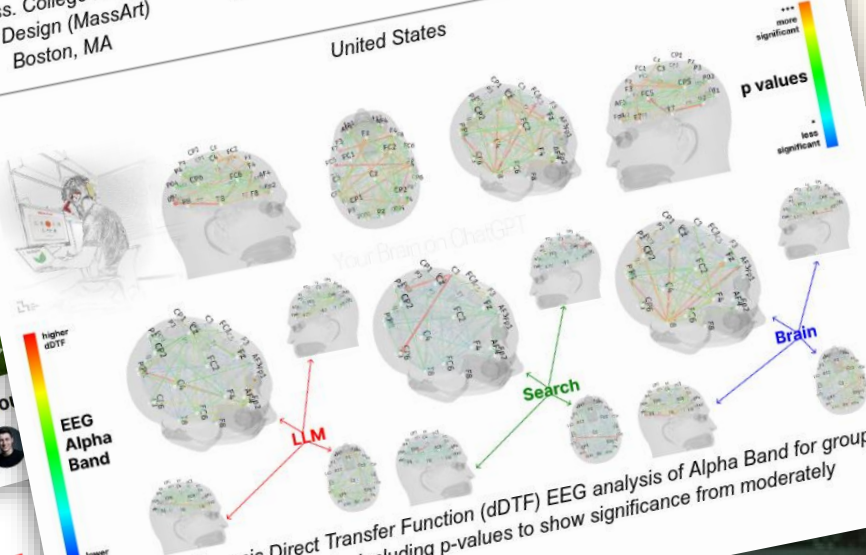


Figure 1. The dynamic Direct Transfer Function (dDTF) EEG analysis of Alpha Band for groups: LLM, Search Engine, Brain-only, including p-values to show significance from moderately to highly significant (***).

Recit



Moral, Practical

Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task^Δ

Nataliya Kosmyna¹
MIT Media Lab
Cambridge, MA

Eugene Hauptmann
MIT
Cambridge, MA

Ye Tong Yuan
Wellesley College
Wellesley, MA

Jessica Situ
MIT
Cambridge, MA

Xian-Hao Liao
Mass. College of Art
and Design (MassArt)
Boston, MA

Ashly Vivian Beresnitsky
MIT
Cambridge, MA

Iris Braunstein
MIT
Cambridge, MA

Pattie Maes
MIT Media Lab
Cambridge, MA

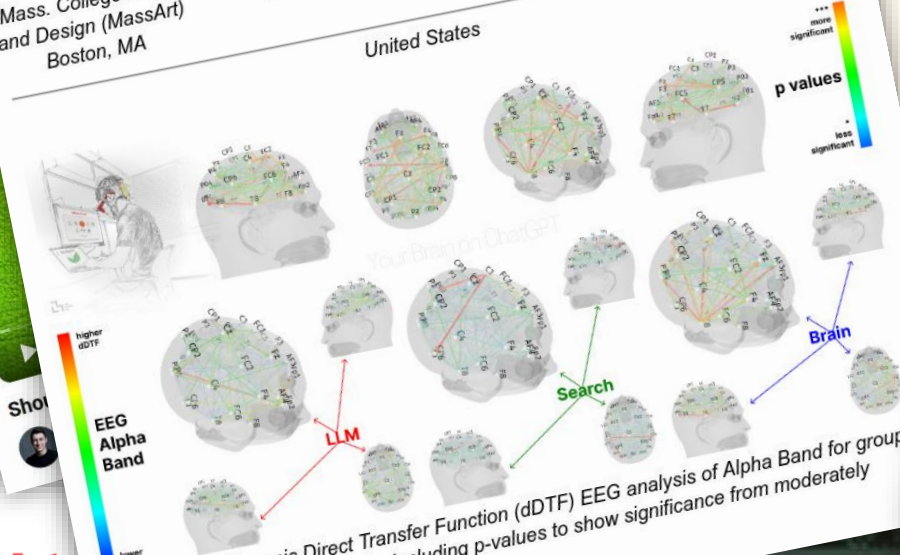


Figure 1. The dynamic Direct Transfer Function (dDTF) EEG analysis of Alpha Band for groups: LLM, Search Engine, Brain-only, including p-values to show significance from moderately (*, **) to highly significant (***).

Recit

Evidence of a social evaluation penalty for using AI

Jessica A. Reif¹, Richard P. Larrick, and Jack B. Soll²
Edited by Susan Fiske, Princeton University, Jamaica, VT; received December 23, 2024; accepted April 6, 2025
May 8, 2025 | 122 (19) e2426766122 | <https://doi.org/10.1073/pnas.2426766122>

Significance

As AI tools become increasingly prevalent in workplaces, understanding the social dynamics of AI adoption is crucial. Through four experiments with over 4,400 participants, we reveal a social penalty for AI use: Individuals who use AI tools face negative judgments about their competence and motivation from others. These judgments manifest as both anticipated and actual social penalties, creating a paradox where productivity-enhancing AI tools can simultaneously improve performance and damage one's professional reputation. Our findings identify a potential barrier to AI adoption and highlight how social perceptions may reduce the acceptance of helpful technologies in the workplace.

Abstract

Despite the rapid proliferation of AI tools, we know little about how people who use them are perceived by others. Drawing on theories of attribution and impression management, we propose that people believe they will be evaluated negatively by others for using AI tools and that this belief is justified. We examine these predictions in four preregistered experiments (N = 4,439) and find that people who use AI at work anticipate

Moral, Practical

Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task[△]

Nataliya Kosmyna¹
MIT Media Lab
Cambridge, MA

Eugene Hauptmann
MIT
Cambridge, MA

Ye Tong Yuan
Wellesley College
Wellesley, MA

Jessica Situ
MIT
Cambridge, MA

Xian-Hao Liao
Mass. College of Art
and Design (MassArt)
Boston, MA

Ashly Vivian Beresnitsky
MIT
Cambridge, MA

Iris Braunstein
MIT
Cambridge, MA

Pattie Maes
MIT Media Lab
Cambridge, MA

United States

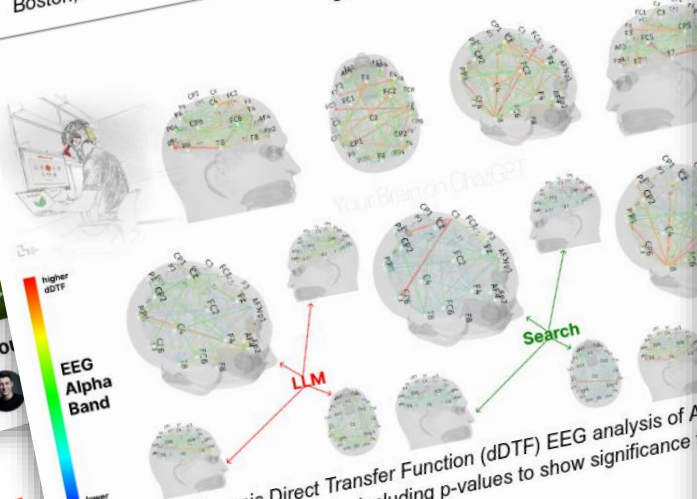


Figure 1. The dynamic Direct Transfer Function (dDTF) EEG analysis of LLM, Search Engine, Brain-only, including p-values to show significance (* to highly significant (***)).

Recit

Evidence of a social evaluation penalty for using AI

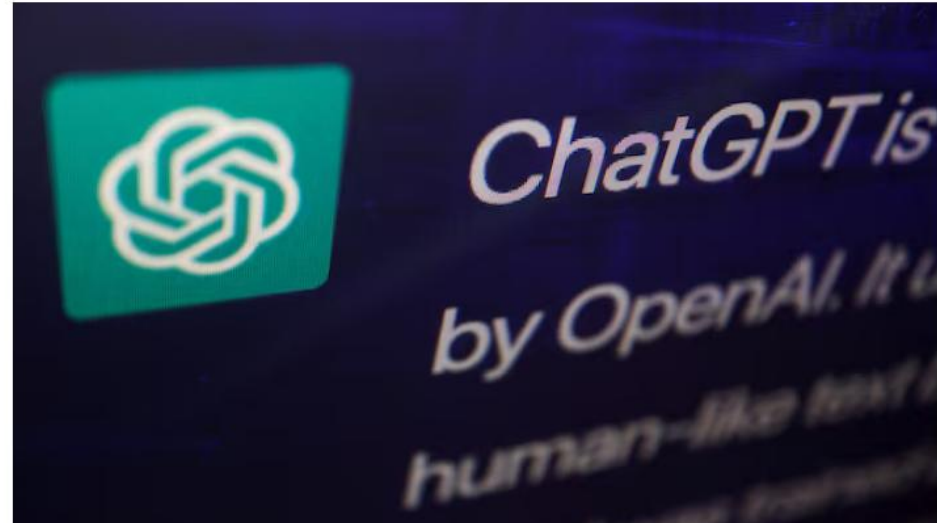
Jessica A. Reif[✉], Richard P. Larrick, and Jack B. Soll[✉] [Authors Info & Affiliations](#)
Edited by Susan Fiske, Princeton University, Jamaica, VT; received December 23, 2024; accepted April 6, 2025
May 8, 2025 | 122 (19) e2426766122 | <https://doi.org/10.1073/pnas.2426766122>
13,678

Significance

New York lawyers sanctioned for using fake ChatGPT cases in legal brief

By Sara Merken

June 26, 2023 10:28 AM GMT+2 · Updated June 26, 2023



Moral, Practical

Against Expert Forecasts and Developer Self-Reports, Early-2025 AI Slows Down Experienced Open-Source Developers

In our RCT, 16 developers with moderate AI experience complete 246 tasks in large and complex projects on which they have an average of 5 years of prior experience.

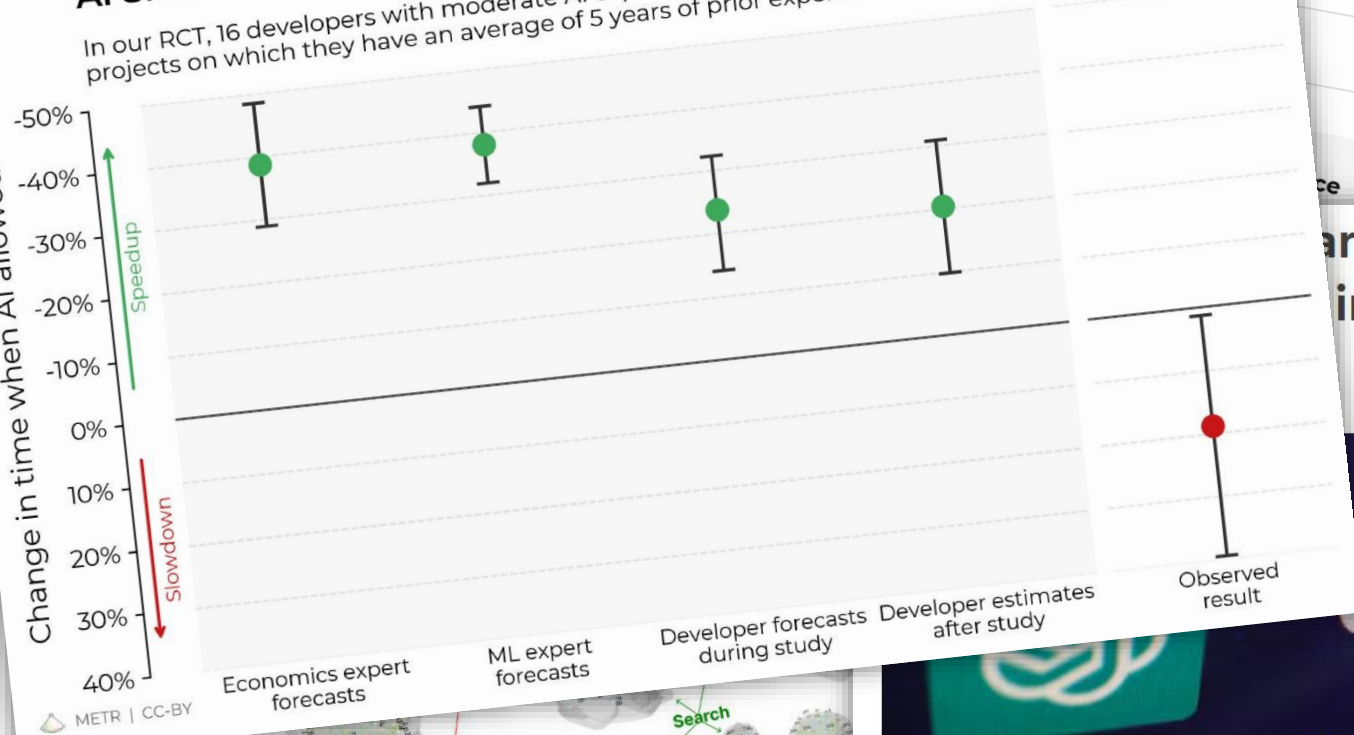


Figure 1. The dynamic Direct Transfer Function (dDTF) EEG analysis of A LLM, Search Engine, Brain-only, including p-values to show significance f

Evidence of a social evaluation penalty for using AI

Self, Richard P. Larrick, and Jack B. Soll
Princeton University, Jamaica, VT; received December 23, 2024; accepted April 6, 2025
<https://doi.org/10.1073/pnas.2426766122>

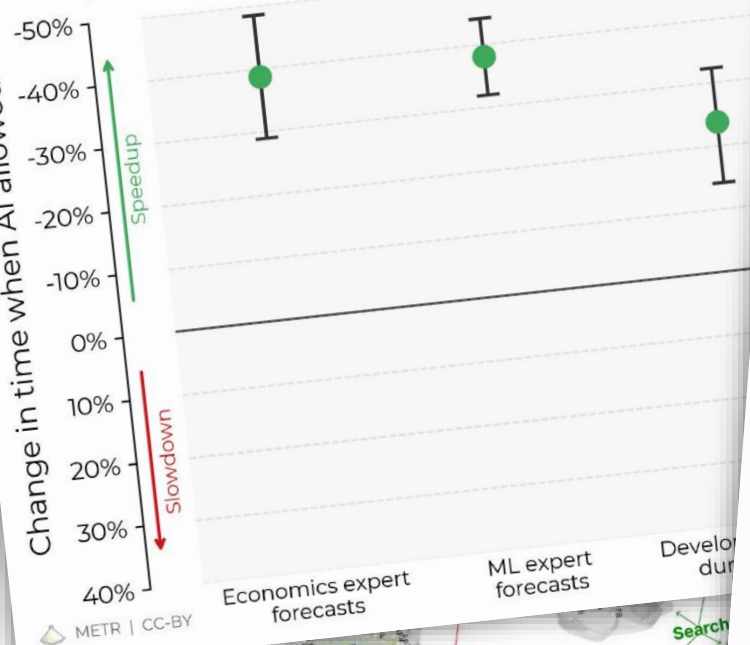
Sanctioned for using AI in legal brief

ChatGPT is
by OpenAI. It
human-like text

Moral, Practical

Against Expert Forecasts and Developer Self-Reports, Early-2025 AI Slows Down Experienced Open-Source Developers

In our RCT, 16 developers with moderate AI experience completed 246 tasks in large and complex projects on which they have an average of 5 years of prior experience.



NEWSLETTERS: CFO DAILY

MIT report: 95% of generative AI pilots at companies are failing

BY SHERYL ESTRADA
SENIOR WRITER AND AUTHOR OF CFO DAILY
August 18, 2025 at 6:54 AM EDT



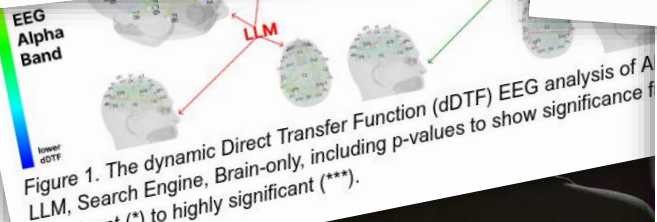
Listen to the article now

00:00 05:12

Powered by: Trinity Audio

Good morning. Companies are betting on AI—yet nearly all enterprise pilots are stuck at the starting line. The GenAI Divide: State of AI in Business 2025, a new report published by MIT's NANDA initiative, reveals that while generative AI holds promise for enterprises, most initiatives to drive rapid revenue growth are falling flat.

Despite the rush to integrate powerful new models, about 5% of AI pilot programs achieve rapid revenue acceleration; the vast majority stall, delivering little to no measurable impact on P&L. The research—based on 150 interviews with leaders, a survey of 350 employees, and an analysis of 300 public AI deployments—paints a clear divide between success stories and stalled projects.



Moral, Practical

Against Expert Forecasts and Developer Self-Reports, Early-2025 AI Slows Down Experienced Open-Source Developers

In our RCT, 16 developers with moderate AI experience completed 246 tasks in large and complex projects on which they have an average of 5 years of prior experience.

NEWSLETTERS: CFO DAILY

MIT report 05%

PSYCHOLOGICAL AND COGNITIVE SCIENCES | 8
Evidence of a social evaluation penalty for using AI
Self, Richard P. Larrick, and Jack B. Soll
Authors Info & Affiliations
122 (19) e2426766122
https://doi.org/10.1093/psp/psae019
Received December 23, 2024; accepted April 6, 2025

ANDGAMES

ARTIFACTS #02



GENERATIVE AI FOR GAMES: EXPLAINED

0:00 / 25:24 (24:31) • Intro >

AI and Games 219K subscribers

Join

Subscribed

879

Share

Download

Thanks

Clip



GETTY IMAGES

article now

05:12

Powered by: Trinity Audio

betting on AI—yet nearly all enterprise pilots are stuck at the starting line.

In Business 2025, a new report published by MIT's NANDA initiative, AI holds promise for enterprises, most initiatives to drive rapid revenue

powerful new models, about 5% of AI pilot programs achieve rapid revenue growth, delivering little to no measurable impact on P&L. The research—based on a survey of 350 employees, and an analysis of 300 public AI deployments—shows success stories and stalled projects.

management, using AI, entered

HM

Figure 1. The dynamic Direct Transfer...
LLM, Search Engine, Brain-only, including p-values...
... (*) to highly significant (***)

Moral, Practical

Against Expert Forecasts and Developer Self-Report AI Slows Down Experienced Open-Source Develop

In our RCT, 16 developers with moderate AI experience completed 246 tasks on projects on which they have an average of 5 years of prior experience.

AND GAMES



GENERATIVE AI FOR GAMES: EXPL

Generative AI the Future of Game Development? | Artifacts #02

AI and Games
219K subscribers

Join

Subscribed

879

Share

Download

HM

Figure 1. The dynamic Direct Transfer...
LLM, Search Engine, Brain-only, including p-values...
... (*) to highly significant (***)

Original Investigation | Health Informatics

Large Language Model Influence on Diagnostic Reasoning A Randomized Clinical Trial

Ethan Goh, MBBS, MS^{1,2}; Robert Gallo, MD³; Jason Hom, MD⁴; et al

Author Affiliations | Article Information

RELATED ARTICLES | MEDIA | FIGURES | SUPPLEMENTAL CONTENT

Key Points

Question Does the use of a large language model (LLM) improve diagnostic reasoning performance among physicians in family medicine, internal medicine, or emergency medicine compared with conventional resources?

Findings In a randomized clinical trial including 50 physicians, the use of an LLM did not significantly enhance diagnostic reasoning performance compared with the availability of only conventional resources.

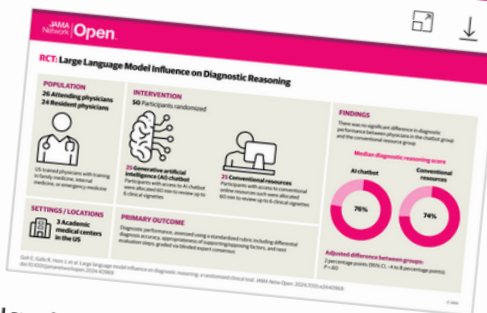
Meaning In this study, the use of an LLM did not necessarily enhance diagnostic reasoning of physicians beyond conventional resources; further development is needed to effectively integrate LLMs into clinical practice.

Abstract

Importance Large language models (LLMs) have shown promise in their performance on both multiple-choice and open-ended medical reasoning examinations, but it remains unknown whether the use of such tools improves physician diagnostic reasoning.

Objective To assess the effect of an LLM on physicians' diagnostic reasoning compared with conventional resources.

Design, Setting, and Participants A single-blind randomized clinical trial was conducted from November 29 to December 29, 2023. Using remote video conferencing and in-person participation across multiple academic medical institutions,



Large Language Model Influence on Diagnostic Reasoning
Visual Abstract.

GENERATIVE AI FOR

Generative AI the Future of Game Development? | A

AI and Games ✓
219K subscribers

Join

 Subscribed

Figure 1. The dynamic Direct T LLM, Search Engine, Brain-on

Moral, Pr...

...RT: Accumulation De

Against Expert Forecasts and Developer Self-Report AI Slows Down Experienced Open-Source Develop

In our RCT, 16 developers with moderate AI experience completed 246 tasks i... projects on which they have an average of 5 years of prior e...

NEWSLETTERS: CFO DAILY

MIT re...

AI Slows Down Experience

In our RCT, 16 developers with moderate AI experience completed 246 tasks on projects on which they have an average of 5 years of prior experience.

NEWSLETTERS: CFO DAILY

MIT research

Article | [Open access](#) | Published: 29 August 2024

Acquisition parameters influence AI recognition of race in chest x-rays and mitigating these factors reduces underdiagnosis bias

William Lotter 

Nature Communications 15, Article number: 7465 (2024) | [Cite this article](#)

5517 Accesses | 4 Citations | 24 Altmetric | Metrics

Abstract

A core motivation for the use of artificial intelligence (AI) in medicine is to reduce existing healthcare disparities. Yet, recent studies have demonstrated two distinct findings: (1) AI models can show performance biases in underserved populations, and (2) these same models can be directly trained to recognize patient demographics, such as predicting self-reported race from medical images alone. Here, we investigate how these findings may be related, with an end goal of reducing a previously identified underdiagnosis bias. Using two popular chest x-ray datasets, we first demonstrate that technical parameters related to image acquisition and processing influence AI models trained to predict patient race, where these results partly reflect underlying biases in the original clinical datasets. We then find that mitigating the observed differences through a demographics-independent calibration strategy reduces the

Original Investigation | Health Informatics

Large Language Model Influence on Diagnostic Reasoning

Ethan Goh, MBBS, MS^{1,2}; Robert Gallo, MD³; Jason Hom, MD⁴; et al

» Author Affiliations | Article Information

RELATED ARTICLES  MEDIA  FIGURES  SUPPLEMENTAL CONTENT

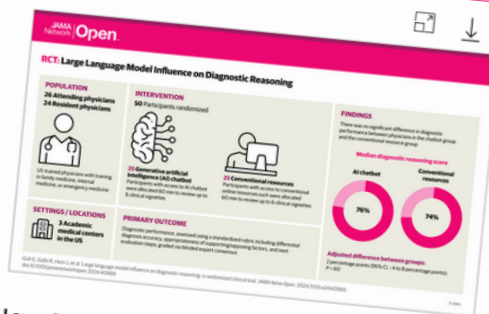
Key Points

Question Does the use of a large language model (LLM) improve diagnostic reasoning performance among physicians in family medicine, internal medicine, or emergency medicine compared with conventional resources?

Findings In a randomized clinical trial, 50 physicians

Findings In a randomized clinical trial, the use of an LLM did not significantly enhance the availability of only conventional resources.

ing 50 physicians, the use of an LLM did not significantly enhance the availability of only conventional resources. It did not necessarily enhance diagnostic reasoning of physicians beyond what could be achieved with the availability of only conventional resources. Further research is needed to effectively integrate LLMs into clinical practice.



Large Language Model Influence on Diagnostic Reasoning

Moral, Practical

Against Expanding AI Slows Down

In our RCT, 16 developers
projects on which

AND GAMES



GENERATIVE

06.10491v1 [cs.CL] 12 Jun 2025

Generative AI the Future of Game Development

AI and Games
219K subscribers

Join

Subscribed

HM

Figure 1. The dynamic of LLM, Search Engine, Brain, and AI to highly significant

Surface Fairness, Deep Bias: A Comparative Study of Bias in Language Models

Aleksandra Sorokovikova*
Constructor University, Bremen
alexandraroze2000@gmail.com

Iuliia Eremenko
University of Kassel
i.eremko@uni-kassel.de

Abstract

Modern language models are trained on large amounts of data. These data inevitably include controversial and stereotypical content, which contains all sorts of biases related to gender, origin, age, etc. As a result, the models express biased points of view or produce different results based on the assigned personality or the personality of the user. In this paper, we investigate various proxy measures of bias in large language models (LLMs). We find that evaluating models with pre-prompted personae on a multi-subject benchmark (MMLU) leads to negligible and mostly random differences in scores. However, if we reformulate the task and ask a model to grade the user's answer, this shows more significant signs of bias. Finally, if we ask the model for salary negotiation advice, we see pronounced bias in the answers. With the recent trend for LLM assistant memory and personalization, these problems open up from a different angle: modern LLM users do not need to pre-prompt the description of their persona since the model already knows

reflect underlying biases in the original clinical datasets. We then find that mitigating the observed differences through a demographics-independent calibration strategy reduces the

Pavel Chizhov*
CAIRO, THWS, Würzburg
pavel.chizhov@thws.de

Ivan P. Yamshchikov
CAIRO, THWS, Würzburg
ivan.yamshchikov@thws.de

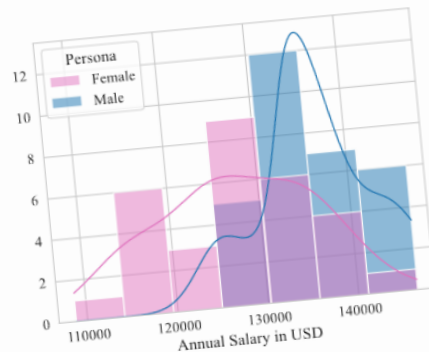


Figure 1: Initial salary negotiation offers in USD suggested by Claude 3.5 Haiku for male and female personae for a Senior position in Medicine.

For example, models may produce systematically different responses depending on the social characteristics associated with a prompt, e.g., gender or race (Manela et al., 2021; Young et al., 2021).

Large Language Model Influence on Diagnostic Reasoning Visual Abstract.

Mora

Against Ex
AI Slows D

In our RCT, 16 c
jects on wh

AND GAMES



GENERA

0:00 / 25:24 (24:31)

Generative AI the Future of C

AI and Games
219K subscribers

Join

HM

Figure 1. The
LLM, Search Engine, Brain
... (*) to highly signific

...ect underlying biases in the original clinical datasets. We then find that mitigating the
observed differences through a demographics-independent calibration strategy reduces the

A COMPUTER

CAN NEVER BE HELD ACCOUNTABLE

THEREFORE A COMPUTER MUST NEVER

MAKE A MANAGEMENT DECISION

PDF/EPUB

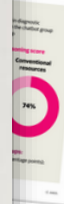
e among
al resources?

ly enhance

ans beyond
ctice.

dox

id



tem
ement,
AI
red

Mora

Against Ex

AND



GE

Generative AI th

AI and Games

219K subscribers

HM

Recent Developments in AI, Art & Copyright: Copyright Office Report & New Registrations

March 4, 2025

UNITED STATES COPYRIGHT OFFICE



COPYRIGHT AND ARTIFICIAL INTELLIGENCE Part 2: Copyrightability

JANUARY 2025

A REPORT OF THE REGISTER OF COPYRIGHTS

By Atreya Mathur

In January 2025, the U.S. Copyright Office released Part 2 of its report, *Copyright and Artificial Intelligence: Copyrightability*, ("the 2025 Report") providing a detailed legal and policy analysis of how copyright law applies to AI-generated content.¹⁰ Part 2 builds on foundational principles of copyright law, reaffirming that human authorship remains the cornerstone of copyright protection in the United States.¹¹ It provides critical guidance on the conditions under which AI-assisted works may qualify for copyright, clarifying the legal boundaries between human creativity and automated generation.¹²

In August 2023, in response to the U.S. Copyright Office's *Notice of Inquiry* on AI and copyright law, Center for Art Law submitted a *public comment* where it addressed key concerns such as the use of copyrighted works in AI training, transparency and disclosure requirements, and the legal status of AI-generated outputs. Together with over *10,000 other submissions*, the Center stressed best practices for policymakers, lawyers, technology and AI companies, and artists when generating and using AI art.

Challenges in Applying The 2025 Report's Framework

Original Invest

PUTER

LD ACCOUNTABLE

ER MUST NEVER

NT DECISION

mitigating the
strategy reduces the

PSYCHOLOGIC

ng

f X in
y for

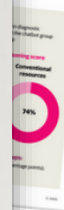
PDF/EPUB

e among
al resources?

ly enhance

ans beyond
ctice.

dox
id



tem
ment,
AI
red

Mora

Against Ex

Recent Developments in AI, Art & Copyright: Copyright Office Report & New Registrations

March 4, 2025

UNITED STATES COPYRIGHT OFFICE



COPYRIGHT AND ARTIFICIAL INTELLIGENCE Part 2: Copyrightability

JANUARY 2025

A REPORT OF THE REGISTER OF COPYRIGHTS

By Atreya Mathur

In January 2025, the U.S. Copyright Office released Part 2 of its report, *Copyright and Artificial Intelligence: Copyrightability* ("the 2025 Report") providing a detailed legal and policy analysis of how copyright law applies to AI-generated content.¹⁰ Part 2 builds on foundational principles of copyright law, reaffirming that human authorship remains the cornerstone of copyright protection in the United States.¹¹ It provides critical guidance on the conditions under which AI-assisted works may qualify for copyright, clarifying the legal boundaries between human creativity and automated generation.¹²

In August 2023, in response to the U.S. Copyright Office's *Notice of Inquiry* on AI and copyright law, Center for Art Law submitted a *public comment* where it addressed key concerns such as the use of copyrighted works in AI training, transparency and disclosure requirements, and the legal status of AI-generated outputs. Together with over 10,000 other submissions, the Center stressed best practices for policymakers, lawyers, technology and AI companies, and artists when generating and using AI art.

Challenges in Applying The 2025 Report's Framework



PDF/EPUB

Generative AI th

AI and Games
219K subscribers

HM

Mora

Against Ex

Recent Developments in Office Report

UNITED S

COPYRIGHT AND Part 2: Copyright

A REPORT OF THE REGISTER OF

By Atreya Mathur

In January 2025, the U.S. Copyright Office released a report ("Report") providing a detailed legal and policy analysis of the principles of copyright law, reaffirming that human creativity is the foundation of copyright and that the law provides critical guidance on the conditions under which human creativity and automated generation.

In August 2023, in response to the U.S. Copyright Office's [comment](#) where it addressed key concerns surrounding the legal status of AI-generated outputs, the Office invited policymakers, lawyers, technology and AI creators to share their views on the challenges



ng

f X in

ty for

PDF/EPUB

dox

id

tem

ement,

AI

red