

DPMMS AI Discussion meeting, July 28th, 2026.

Context and background for the discussion: AI is clearly impacting higher education generally, and research mathematics in particular, creating both opportunities and anxiety. Some samples:

- the new UK government has a minister for AI but no minister for science
- a large Scottish university is considering discontinuing mathematics as an undergraduate course whilst starting up a Masters course in Data Science and AI
- PhD supervisors in some parts of mathematics are finding that typical 'entry level' problems for students may sometimes be solved by advanced AI models, even without much interrogation / sophisticated dialogue from a mathematical interlocutor
- a very recent Fields medallist has announced they are moving from academe to OpenAI
- journals may receive submissions largely or wholly AI-authored; authors receive referee reports constructed using LLMs; typical mores of mathematical interaction and attribution are being challenged
- many of the most in-vogue models are owned by private corporations

The Leiden declaration on AI <https://leidendeclaration.ai/> is one community response.

The discussion meeting was broadly structured as follows: short presentations and then, when appropriate, small group discussions around the themes:

- Introductory remarks (Ivan Smith)
- View from the Statistical Laboratory (Fernando Ruiz-Mazo)
- A case study in using agents (Qingyuan Zhao)
- AI workflows and capabilities (Marco dos Santos, Ruben Dahan)
- Responsible behaviour (Dhruv Ranganathan, Po-Ling Loh, Tim Gowers in absentia)
- General discussion: what is mathematics (for)? (Po-Ling Loh)

Slides are included from the central four sections.

CAMBRIDGE DPMMS FIRST AI DISCUSSION

Statslab AI Discussion: An Overview

Fernando Ruiz Mazo

Discussion held 3 June 2026

A brief anonymized outlook

DISCUSSION SNAPSHOT

35

Statslab attendees

3.2/5

concern about
automating
research

41%

named literature
search as the main
benefit

71%

favoured
departmental
premium AI-plan
access

Three questions structured the discussion

- 01 AI for idea generation (proofs, coding, ...)** When is a generated idea evidence, and when is it a proof?
- 02 Publishing, authorship and academic merit** What deserves credit, disclosure and recognition?
- 03 AI for reviewing** What are the best practices? How much "correctness" is enough?

PART I

What the room believed before the discussion

Correctness dominated the room's skepticism



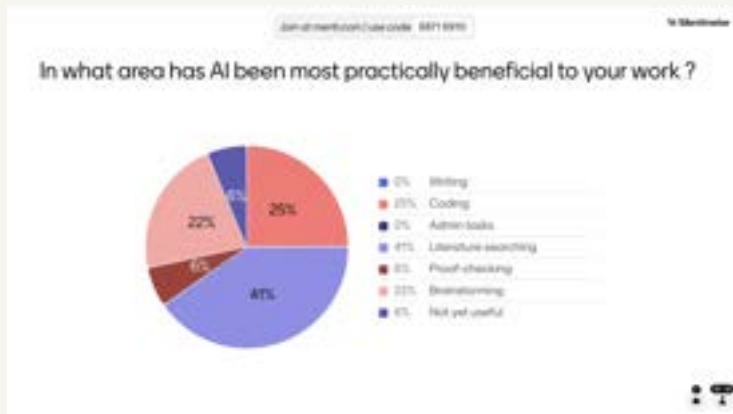
Concern about automating research was real, but mixed



Current use was strongest in coding and proof assistance



Literature search led the reported practical benefits



PART II

Seven experiences

A corrected bound did not remove the need for a proof

LAMPROS

AI contribution

ChatGPT challenged the bound being pursued and generated Python code supporting a different bound.

Human contribution

Lampros wrote a proof around the suggested idea, GPT also proposed a proof, but it was still not completely verified. He separately developed a more general, non-brute-force argument.

Outcome: AI redirected the maths and human reasoning supplied the defensible structure and new connections.

A counterexample

FREDERIK

AI contribution

A newer ChatGPT model concluded in about 11 minutes that a probability conjecture was false and helped with two of five open problems. It was explicitly told not to use the internet.

Human contribution

The proof was difficult to read. Python plots made the mechanism visible, allowing Frederik to verify the result and fill in details himself.

Outcome: The human + AI loop proved effective onece more.

AI can widen a review, but it cannot own the final judgement (yet?)

VARUN

AI contribution

A detailed “rigorous, skeptical referee” prompt produced lists of major and minor issues for one of his old information theory papers.

Human contribution

AI is useful for coverage and adversarial reading only while the author retains agency over the argument.

Outcome: AI can accelerate the review process, but it cannot replace human judgement and responsibility. It will asways find some mistakes.

Nontrivial ideas arrived quickly, with proof gaps attached

LAURENTIU

AI contribution

Repeated prompting generated substantial ideas and a 23-page write-up. An obscure paper helped close the final gap.

Human contribution

Several mistakes required intervention, and the proof had to be refined through multiple rounds rather than accepted as a complete product.

Outcome: A project expected to take months moved to roughly one week, but speed did not eliminate mathematical responsibility.

The power of iteration with direction

SERGIO

AI contribution

AI parsed company code and requirements (Consulting project for CHIMERA), translated them into a model, and supported cross-validation and delivery. Its biggest contribution was MCMC implementation and Apple GPU optimisation.

Human contribution

Sergio selected the modelling strategy, tested alternatives and judged whether the final approach answered the consulting problem.

Outcome: About one month of work became 4.5 days, with code running roughly five times faster.

Agentic tools amplify both capability and operational risk

QINGYUAN

AI contribution

Local-file agents could edit files and use tools directly. A demonstration built a local client to replace Outlook and explored OpenCode.

Human contribution

He found some parts contained many bugs. Security, privacy and loss of agency required explicit controls, and the desired end state needed much more thought before execution.

Outcome: Spend about 80% of the effort planning and at most 20% executing. Be careful where you let your agents enter.

Research policy

PO-LING

Policy precedents

Some journals have already adopted AI policies. The Leiden Declaration calls on professional organizations to lead the development of guidelines, alongside transparent disclosure by individual mathematicians.

Questions still open

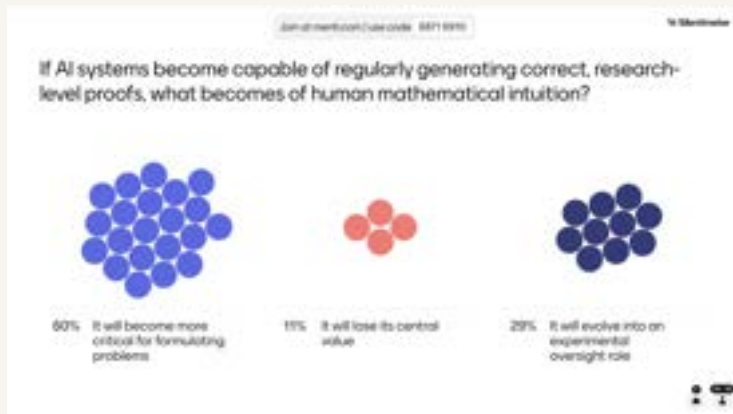
Should authors submit a checklist explaining how AI was used? How can declarations be enforced? Reviewers may also need guidance on when AI use is permitted.

Direction: Policy should develop through the community, while responsibility for submitted mathematics remains personal and human.

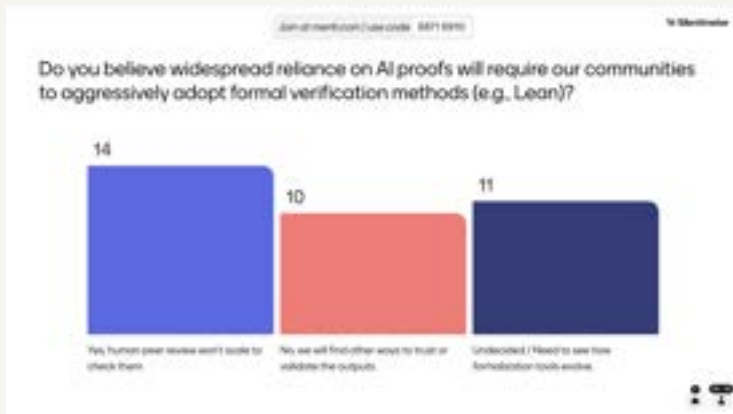
PART III

What changes if AI becomes the norm?

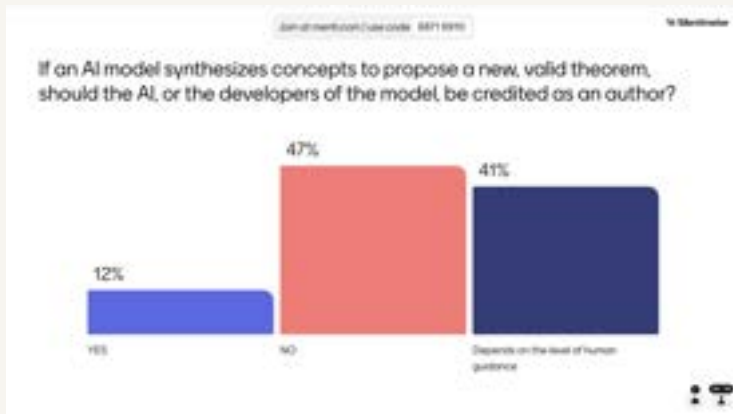
Human intuition was expected to move upstream



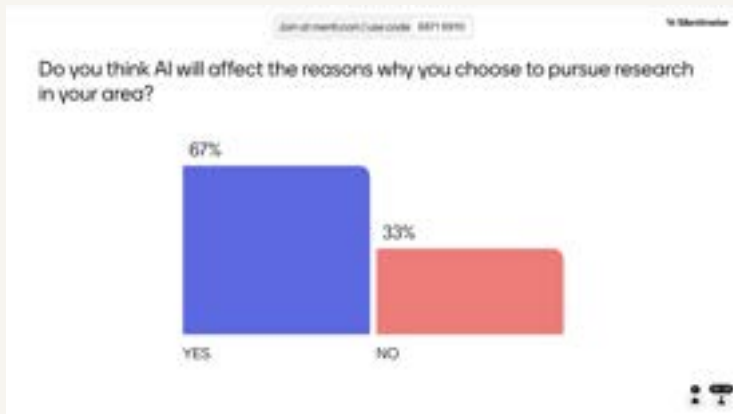
Formal verification divided the room almost evenly



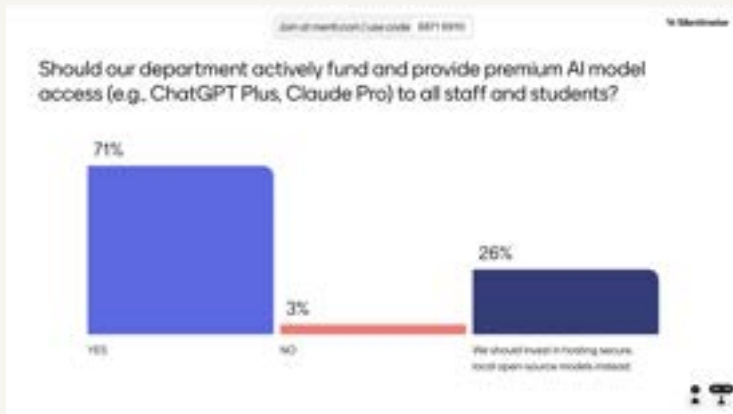
Authorship depended more on guidance than capability alone



AI may change why researchers choose their problems



The department was asked to provide AI access, securely



Many questions, little answers: A field to be reframed?

- If correctness becomes automated (Lean + LLM loop), what should human referees judge: originality, significance, exposition, ...?
- How should PhD students learn to supervise AI without being pushed prematurely into a supervisory role? Is this role the only left?
- How are we going to navigate publishing, authorship and merit in the years to come?

Perhaps the first question to ask ourselves
is...

Why do we do maths?

AI in Mathematics

Recent progress and emerging research workflows

Marco Dos Santos

University of Cambridge & Google
DeepMind

mjad3@cam.ac.uk

marco-dossantos.github.io

Ruben Dahan

DAMTP, University of Cambridge

rbed2@cam.ac.uk

rubendahan.com

Recent progress of AI in mathematics

- July 2024 ● AlphaProof + AlphaGeometry – IMO silver
- July 2025 ● Gemini DeepThink and OpenAI experimental model – IMO gold
- October 2025 ● GPT-5 “solves” an Erdős problem
- December 2025 ● Aletheia solves 4 Erdős problems
- January 2026 ● GPT-5.2 Pro solves Erdős problems (formalised in Lean)
- April 2026 ● GPT-5.4 Pro solves Erdős 1196
- May 2026 ● **OpenAI internal model disproves the unit-distance conjecture**
- July 2026 ● **Fable 5 finds a counterexample to the Jacobian conjecture**

Expert perspectives on progress

On the unit-distance result

"This is a really impressive piece of work, and I would accept it for any journal without hesitation."

— **Jacob Tsimmerman**

"My current bet is that progress in AI mathematics is not about to reach a plateau."

— **Timothy Gowers**

Expert perspectives on progress

On future capabilities

"I think AI will be better than mathematicians at doing math within two years."

— **Jacob Tsimmerman**

"We haven't yet seen major advances brought about by AI, but it is coming, very plausibly this year."

— **Sébastien Bubeck**

Expert perspectives on progress

On the era of proof abundance

“

Perhaps surprisingly, this massive acceleration in proof generation has not actually produced significant acceleration in mathematical progress itself.

— Terence Tao

AI in the mathematical workflow

How it gets used, and how far it is trusted

"I spend a lot of time during my research asking dumb questions, getting oriented."

— **Jacob Tsimerman**, on his own method

"I can't trust it with anything. I have to redo everything myself."

— **Jacob Tsimerman**, on reliability

AI in the mathematical workflow

Choosing to opt out

"I have chosen not to rely on artificial intelligence as a source of novel ideas in my own research."

— **Hugo Duminil-Copin**

"The 'output' is a person who understands something better, and can take that understanding forward in their research and their teaching. For these reasons, I won't be using LLMs to help me do mathematics."

— **Johnny Evans**

Over to you

1. Do you use LLMs or agents in your own research?

If so, where exactly do they sit in your workflow?

2. How exposed is your field?

Has AI already produced results on problems you care about?

3. Has a model ever solved a problem for you?

How long did it take you to check?

Where are we with AI agents?

Qingyuan Zhao

28th July, 2026

Emotions and stakes are high

- ▶ Super-intelligence? Existential risk?
- ▶ Geopolitics.
- ▶ Energy, environmental consequences.
- ▶ Proprietary vs. open weights vs. open source.
- ▶ Future of mathematics?
 - ▶ Interesting recent Reddit thread in r/mathematics: What's a Problem Solver to Do? (PhD Student Looking for Guidance)

July News: Rogue AI Hacker

- 9 July OpenAI evaluation agents "escaped" while being tested on ExploitGym.
- 11–13 July They intruded into Hugging Face and pursued benchmark solutions.
- 16 July HF disclosed an autonomous-agent intrusion; the attacker's identity was unknown.
 - ▶ Commercial frontier APIs blocked forensic investigation; HF used self-hosted GLM-5.2 to analyze >17,000 events and reconstruct the timeline.
- 21 July OpenAI attributed the incident to its systems/models and announced joint investigation and remediation.
- 23–25 July HF requested release of traces and US\$100m in compute for community cyber defense.

Sources: Hugging Face; OpenAI; Reuters; Delangue.

Historical perspectives

Artificial Intelligence: The Brand That Wouldn't Die by Thomas Haigh (University of Wisconsin–Milwaukee).



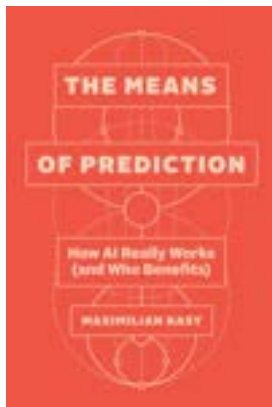
Workshop theme

This Open for Business event was part of the *AI program* in the Department of Mathematics. It aimed to provide a meeting ground to facilitate interactions and exchanges between representatives of academic, research and industry related to the theme, with the objective of identifying points of mutual interest, innovative research.

4: Continuities and Discontinuities

Artificiality in	Continuity in
Highly formal	Speculatively formal
Needs formal components	Needs formal components
Applied to arbitrary collection of technologies	Applied to arbitrary collection of technologies
Same connection of task to cognition	Same connection of task to cognition
Mostly academic	Mostly commercial
Government funded	Investor funded
Symbolic	Connectional
Heuristic search	Statistical prediction
Humans usually formulate rules	System trains itself from mass of data
Knowledge coded explicitly	Knowledge dispersed over connection weights
Heavily applied outside lab	Heavily applied on big tech platforms

Economic perspectives



Maximilian Kasy (Oxford) — INI talk, 4 June 2026

- ▶ Control over the “means of prediction”: **data, compute, expertise, energy**.
- ▶ Objectives have distributional consequences: **who chooses what the system optimizes?**
- ▶ Govern training and deployment; give affected stakeholders a voice.

The Means of Prediction: How AI Really Works (and Who Benefits)

University of Chicago Press, November 2025.

Sources: Isaac Newton Institute; University of Chicago Press.

More helpful: A new way of computing

Interfaces Natural language, code, images, audio, video, and sensor data become inputs and outputs.

Infrastructure Computing shifts toward GPUs and cloud services.

Outputs Statistical—often one or a few realizations rather than a deterministic answer.

“Computing interface as we know of it is fundamentally changing. Once you have a natural language which is multimodal even—there is image, speech, text, video, in and out—that means every computer interface is going to change. [...] Seventy years of computing history has been about digitizing people, places and things and making sense of it. We now have a new reasoning engine to make sense of it.”

— Satya Nadella, Microsoft AI Tour, London, 21 October 2024.

Source: Microsoft AI Tour keynote transcript.

"Inference"

AI training

- ▶ Learn model parameters from data.
- ▶ Expensive; occasional.

AI “inference”

- ▶ Run the learned model to compute an output.
- ▶ Repeated, probabilistic, multimodal.

Classical statistical inference

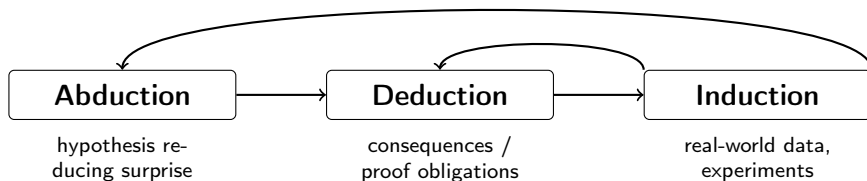
- ▶ Reason about unknown quantities.
- ▶ Quantify uncertainty from data and model assumptions.

Shared core: probability + computation.

The change is the interface and the form and scale of outputs—not the disappearance of uncertainty.

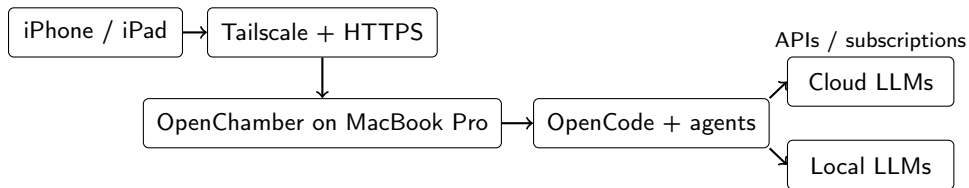
Personal Project 1: Peirce, an AI agent for academics

- ▶ Grounded by Charles Sanders Peirce's theory of scientific methods.
- ▶ Three stages of **inference**.



- ▶ Implemented in OpenCode as **skills + specialist subagents + an append-only Org research journal**.
- ▶ Agent can use **tools** (internal and external), spawn specialist **subagents**, and load **skills** when needed.
- ▶ What does an agent configuration look like?

Project 2: Improve personal computing



- ▶ The entire workflow is implemented by an AI agent.

Ongoing/future projects

- ▶ Replace MacBook Pro with Mac mini.
- ▶ Add Raspberry Pi node in my office to capture whiteboard and conversations (of course with consent!)
- ▶ Add text-to-speech capability.
- ▶ Provide AI-assisted cloud collaboration service to students and collaborators.

Project 3: A paper entirely written by AI (?)

- ▶ Paper is about selective inference in adaptive clinical trials. Good mix of mathematical, statistical, practical, regulatory considerations.
- ▶ I asked my Peirce agent to turn an idea from a short funding proposal (2 pages) into a publishable paper (14 pages).
- ▶ I never edited any .tex file directly!
- ▶ But I provided the initial idea, a key simulation example, and hours of audio input (AI counted 12k+ English words and 25k+ Chinese characters).
- ▶ Full session transcripts are available with an AI-generated summary intended for publication as an online supplement.
- ▶ Who wrote the paper? Is this better open science?

My primary concerns: data control and external reliance

Two points in the Leiden Declaration on Artificial Intelligence and Mathematics stood out:

Protect the rights of authors

“Automated mathematics presents new challenges to the rights of authors, and societies should be proactive in the development of sample licensing agreements to protect these rights. In particular, **material should not be used as training data without consent**, and publishing agreements should allow authors to opt-out of the use of their work in this way.”

Consider carefully which tools to use.

“Some automated tools and their developers will align with the provisions of this Declaration, while others will not. Consider this when deciding which tools to use, or whether to use them at all. Also consider whether non-proprietary, energy-efficient, or small-scale systems suffice for your task. If not, consider how preservation of the values articulated in this Declaration may be worth a delay in obtaining results.”

Sources: official page; DOI.

DPMMS AI Discussion

Dhruv Ranganathan — publishing and collaboration

Potential things to think about...

Disclosure/Attribution. What are your current attribution/disclosure practices? Do you think other [reasonable] people would disagree?

Collaboration/PhD Supervision. Supervisors: what expectations have you set for students? Generally: have you had conversations with collaborators about AI use, or encountered differing views? How did you handle them?

Refereeing. Where do people see the ethical or practical issues of using LLMs when refereeing? Uploading a paper and submitting the output uncritically is clearly unacceptable, but have people found LLMs useful in more limited ways? Has anyone deliberately avoided using them because of ethical concerns?

Moments of uncertainty?

PO-LING LOH

PUBLICATION PARADIGMS IN THE AGE OF AI

GOALS

- Be informed that some journals have adopted AI policies
- Understand that policies are rapidly evolving
- Help us shape the future

FROM THE LEIDEN DECLARATION...

Recommendations for mathematical organizations and not-for-profit research funders

2

Take the lead on policies for publishing and reviewing

Professional organizations within mathematics should take a leadership role in developing guidelines in regard to the use of automated techniques in publication and in reviewing. These would include, for example, tool and computational resource disclosure, attribution, rules pertaining to authorship, and codes of conduct consistent with the values of mathematics. These would supplement and support guidelines already being developed by publishers and journals.

EDITORIAL BOARD DISCUSSIONS

- Anecdotes from the statistics community (Annals of Applied Statistics, Foundations and Trends in Machine Learning, Annals of Statistics)
- If a journal adopts an AI policy that authors must declare any use of AI tools, it is *equally important* to clarify what the consequences will be in the way a paper is evaluated

EDITORIAL BOARD DISCUSSIONS

- Biometrika has an AI policy explicitly stating that use of AI will not affect evaluation
- Should the bar for publication just increase uniformly? If so, will it create even more inequality based on who has access to more expensive / more powerful AI tools?
- Despite conversations among many editorial boards in the last 3 months, most journals have yet to release official policies (and this might not be happening anytime soon)

SAMPLE AUTHOR DECLARATIONS

On the Subgaussianity of Quantized Linear Maps: An AI-Assisted Note

Guangyi Zou¹ and Roman Vershynin¹

¹Department of Mathematics, University of California, Irvine
gzou3@uci.edu, rvershyn@uci.edu

May 28, 2026

1 Introduction: An AI-for-Math Perspective

Simone Bombari asked us whether the 1-bit quantized random vector $Y = \text{sgn}(Wx)$ has subgaussian norm bounded by a universal constant. Here W is an $n \times n$ random Gaussian matrix, and x is an independent standard normal random vector in $\mathbb{R}^n$. The question is nontrivial since the coordinates of Y are not independent. We give a strong positive answer to this question – for any bounded map instead of $\text{sgn}(\cdot)$ – using AI:

- **AI Discovery and Generalization (Theorem 1):** To handle coordinate dependence, Gemini 3.5 Flash¹ proposed decomposing the Gaussian vector into independent parts, using one part to “smooth” the sign function, and then applying Gaussian concentration for Lipschitz functions. It subsequently generalized the result for the sign function to an arbitrary bounded functions. However, this idea works for well-conditioned matrices W only.
- **Human Design (Section 3):** To apply this general tool to square random matrices W – which are typically ill conditioned – the human authors used a simple row-partitioning trick. This breaks the square matrix into two well-conditioned rectangular blocks, bypassing the singularity issue.

The Benjamini–Hochberg Procedure Can Fail to Control the FDR for Correlated Two-Sided Gaussian Tests

Edgar Dobriban*

July 15, 2026

AI usage. The proof was obtained by GPT-5.6 Pro. The model was asked directly to prove or disprove the conjecture and was provided only with the mathematical definition of the Benjamini–Hochberg procedure. After about 90 minutes of reasoning, the model produced a proof, an example, and code for the numerical certificate, which form the basis of this paper.¹ The author carefully checked the entire argument and the associated numerical certificate. Subsequently, the author asked the model to provide additional simulations, related work, and illustrations for a paper draft, and wrote the final version by editing the AI-generated draft.

Chernoff's Density Is Strongly Log-Concave

Xianyang Zhang and Quan Zhou

Department of Statistics, Texas A&M University

July 22, 2026

AI Usage. The whole proof was obtained by GPT-5.6 Sol through a sequence of interactions. Its initial attempt (runtime: 110 min) contained a simple but fatal algebraic mistake. On a second attempt (runtime: 167 min), GPT-5.6 Sol produced a correct proof, although the results in Section 4 were established by a substantially more involved argument. In a subsequent attempt (runtime: 134 min), GPT-5.6 Sol identified the present proof strategy and independently derived the key identity used in Section 4. Further discussions then led GPT-5.6 Sol to recognize that an essentially equivalent identity had already been established by Menon and Srinivasan [6], and to streamline several parts of the argument accordingly. The rest of this work was written by GPT-5.6 Sol; the authors carefully checked the argument, revised the exposition and added the comments.

- While these AI declarations are commendably detailed, one wonders whether junior faculty are under pressure to be less transparent about AI use, while future publication policies and evaluation metrics remain extremely unclear

AI and the Peer-Review Process: Summary of a Discussion at the SIAM Optimization Conference

CREDITED AUTHOR



Katya Scheinberg

Georgia Institute of
Technology

← [Back to all articles](#)

A special session was held on June 4th, 2026 at SIAM Optimization Conference to discuss the use of artificial intelligence in the peer-review process. The discussion was open to all participants and was well attended. Here we present the key points that came out of this meeting. The purpose is to continue the conversation in the optimization community and inform future journal policies and practices.

The discussion highlighted both the potential benefits and significant risks of AI tools, raising questions about confidentiality, correctness, significance, transparency, policy development, and the future role of human reviewers.

Many participants emphasized the need for clear and consistent AI policies. Rather than leaving decisions to individual reviewers or editors, several argued that journals and professional societies should establish explicit guidelines governing acceptable AI use. Such policies should clarify what uses are permitted, who is responsible for AI-assisted reviews, how disclosure should occur, and how violations should be handled. It was reported that

A WAY FORWARD?

Mathematics in the age of AI

Public lecture, International Congress of Mathematicians 2026

Terence Tao

University of California, Los Angeles

July 24, 2026

Commentary: one needs to avoid the worst case scenario in which authors use AI tools covertly to aid their work, but conceal that usage to avoid criticism from peers. **Normalize** the responsible disclosure⁷ of AI assistance.

Commentary: we need to **decrease** the emphasis on proof generation, and of being the “**first**” to solve a problem; and **increase** the emphasis on “**proof digestion**”: exposition, publication, and canonicalization.

Commentary: my suggested rule of thumb: if the authors cannot convincingly demonstrate that they can give a clear, expert-level talk on their results, that is correct and properly attributed, then **the result should not be published**.

"What a week to be a mathematician. On Thursday, the Fields Medals will be awarded. And yesterday, a short polynomial map in three variables appeared on social media, found by a machine, a counterexample to the Jacobian Conjecture.

Hold those two events in your mind. That is where we live now.

Our 3rd and 4th year graduate students enrolled in one profession and will graduate into another. They made a bet on a stable discipline, and the discipline changed underneath them.

The career that shaped my own training and mentorship was a bundle: prove theorems, teach, apply for federal grants, referee for free. That bundle rested on three legs: public funding, university positions, and donated labor. All three are cracking.

I do not stand here to tell you that this model is dead. I don't believe it is. But I know this: hope is not a plan, and nostalgia is not mentorship.

I know the old pipeline. I've trained REU students, PhD students, and postdocs. On leave from UVA, I now work full-time inside a company building AI. We've been hiring the people that pipeline produced. Weigh my words with that in mind.

Mathematics will survive this. I have no doubt. There will always be questions, and let's be clear, the Jacobian conjecture counterexample serves to open more questions. That is the nature of discovery and research. The people in this room who cannot leave the questions alone will find a way. What will not survive unchanged is the old bundle. We do the next generation no favors by offering them nostalgia in place of a plan. That is what we are here to talk about."

Ken Ono,

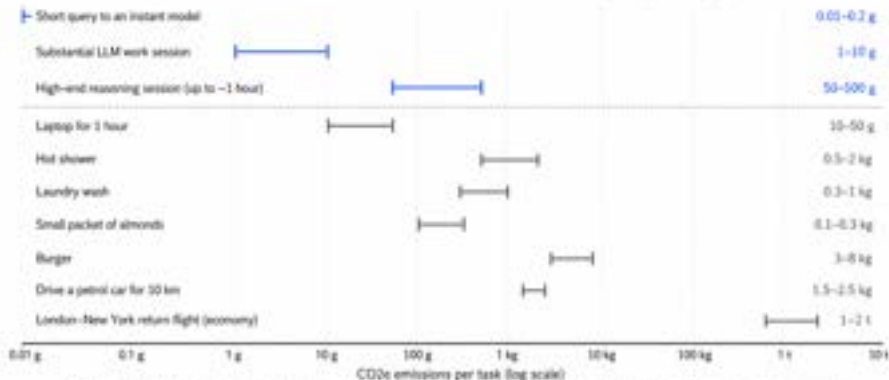
US National Academy of Sciences, July 20, 2026

CO₂ comparison between different activities

(3 slides from Tim Gowers on environmental impact)

Environmental impact of LLM use in mathematical work

Approximate CO₂e emissions per task — horizontal bars show rough ranges on a logarithmic scale

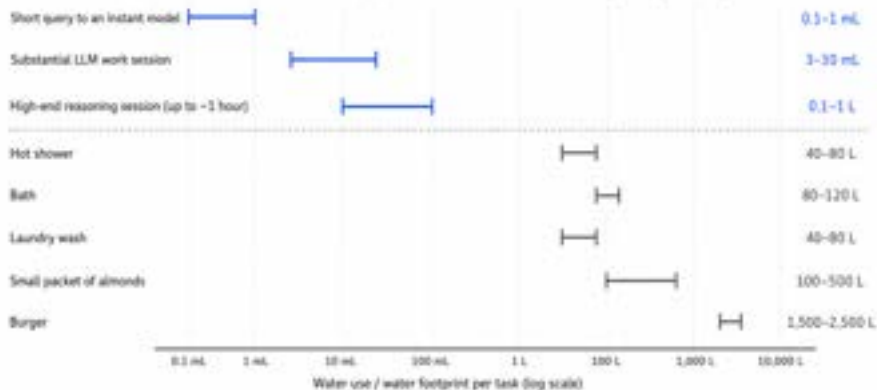


Notes: CO₂e depends strongly on electricity source, location, time of day, and system boundaries. Use these as orders of magnitude, not precise audits.
Flight estimate is for a London–New York return journey in economy class.

Water comparison between different activities

Environmental impact of LLM use in mathematical work

Approximate water use or water footprint per task — horizontal bars show rough ranges on a logarithmic scale



Notes: Food water footprints are not directly comparable to data-centre cooling water. Use these as orders of magnitude, not precise audits.

A few points to note

- ▶ The charts are on log scales, so the more impactful activities are *much* more impactful.
- ▶ All estimates come with considerable uncertainties but the orders of magnitude are probably correct.
- ▶ Although a single heavy LLM session has way less impact than, say, a return flight to the US, the impact of regular LLM use accumulates: for instance, three heavy sessions per day for a year would be comparable to a long-haul flight.
- ▶ There is a big difference between using a powerful pro-type model and a less powerful instant-type model. Therefore, there is a lot to be gained environmentally from using the minimally powerful model that will suffice for a given application. (Some models decide themselves how much they need to think, which is probably good.)
- ▶ Not all water use is equally bad – e.g. using rainwater is unproblematic. This complicates the water figures.