

WHERE I STAND ON AI

VLADIMIR LAZIĆ

At the end of July 2026 I posted on my website a set of principles by which I abide with respect to the use of AI in mathematics [1]. I believed the motivation for such a step was clear, as well as the implicit reference to certain recent events in mathematics. It has, however, become clear that neither the motivation nor the recent events which prompted me to post these principles are known enough. In this document I will clarify both of these points, but the scope of this text goes beyond that original purpose.

Events in May–July 2026. Prior to May 2026 I had mostly ignored the developments in the use of AI in mathematics, as they used to touch problems which I considered irrelevant for the progress of mathematics in general (such as a plethora of activities related to the trivialities of Erdős problems). Several simultaneous events, starting from mid-May 2026, changed completely my views on the imminent dangers of the use of AI in mathematics, and by extension on the society at large.

Around that time, several colleagues started posting numerous papers on the arXiv, under the premise “AI generated, human verified”. As for many other things in this saga, the main mystery for me is *why* someone, who is otherwise an admirable mathematician, would make such a step. Others followed, so that each new day brings several such papers from various mathematicians of various degrees of excellence. As I argue below, this trend is very difficult to explain since it amounts to self-sabotage.

Simultaneously, two Fields medallists, as well as several other mathematicians whose common thread is active presence on social media or the blogosphere, started actively agitating for the increased use of AI in mathematics; a curated selection is [2], [3], [4].

The final straw, not only for me but for many others, was the decision of a third Fields medallist to join an AI company immediately after being awarded the medal in July. The act itself was not enough: since then, he has given several interviews purporting that mathematical research, as we know it, is dead [5], or proclaimed, euphemistically and with a curious smile, that the future of young people in mathematics will, as he puts it, be upended [6]. Not only is this a disgraceful attempt to kill off his own discipline: never has an award been devalued so quickly.

Ethical and environmental issues. There are many contributions online detailing the general ethical issues with the use of AI, as well as disastrous

Mid-August 2026.

consequences for our environment. If one knows them and still does not care, then one is beyond help. As these are well known, I concentrate here on issues which are somewhat less often discussed. There might be some overlap with Section 1 of [7].

Acts of a human. What is the point of living? I do not know, but I am certain that giving up on activities that our species has been doing for thousands of years (art, literature, science) is not what living should be. I refuse to use AI in mathematics (or anywhere else) because giving up *any* skill diminishes what it means to be human. Why would you cut your arm so that your neighbour can grow a third leg?

Good research practice. I refer to the guidelines on good research practice published by the German Research Foundation, which is an authoritative source [8]. In recent months, the use of AI in research has started eroding the most fundamental values of good research practice: here I refer in particular to Guidelines 2, 9, 11, 12 and 14 in [8].

I discuss in detail only one disturbing new trend, which destroys the very foundations of good research practice. That recent trend is an abundance of papers on arXiv which are completely AI generated, but carry “authorship” of humans, usually under the motto “AI generated, human verified”. The strategy is this: a human types a prompt in one (or more) AI agents to solve or disprove one (or more) conjectures in mathematics. After many hours of AI activity (often spread between several AI agents by different companies in order to maximise chances of success and minimise potential for false proofs), a first draft is created. The human “author” is often not an expert in the methods of the draft: in order to create some form of insurance that the mathematics is sound, another AI agent is used to auto-formalise the draft in a proof assistant, creating sometimes thousands of lines of code that no human would quickly be able to verify. After the proof assistant “confirms” that the proof is sound, another AI agent is used to write the draft up in $\text{\TeX}$. At this stage the human steps in, polishing the new version, adding some references in order to make it appear human as if they know what they are talking about, and posting it to the arXiv.

This is usually accompanied by comments that the work was obtained in “conversations with AI”, that there was some “human input”, that “due to the limitation of generative AI it is possible that some literature was missed”, that there was some/extensive “human verification and polishing”, that “the correctness of the paper is the sole responsibility of the authors”. Many references to the “author” are sprinkled throughout, as if referring to oneself as an “author” makes you one. And probably that is the point: nowadays, if you claim you are an author, then you surely must be one.

The explanations of Guideline 14 from [8] include that taking part in “the drafting of the manuscript” could warrant authorship of the manuscript. The current trend is to interpret this guideline in such a way that taking part in drafting of a manuscript is the *sole reason* to warrant authorship.

Say you are a professor who gives a problem to a PhD student. After struggling with the problem for some months, the student comes back with a write-up of a solution. Say you thank the student for the hard work, make aesthetic changes to the manuscript and post it to the arXiv as the sole author, with comments that the work is “student generated, professor verified”, that the work was obtained in “conversations with the student”, that there was some “professor input”, that “due to the limitation of the student it is possible that some literature was missed”, that there was some/extensive “professor verification and polishing”, that “the correctness of the paper is the sole responsibility of the professor”. What is the difference?

The difference, I already hear some say, is that a student is human, and AI is a machine; a student can complain about good research practice, and the machine cannot or does not “want to”. Ironically, the same people also say that their papers are obtained in “conversations with AI”, and happily engage in discussions on whether AI is “sentient”.

I repeat one of my principles: You are not an author of a paper to which your only contribution was to type a prompt in an LLM. If you put your name on such a paper, you are committing an ultimate violation of good scientific practice. You are not a mathematician unless you create new theories or prove new results: prompting a machine to do the work for you does not make you into a mathematician. Asking a machine to do the work for you does not warrant your authorship of a paper.

The futures. Mathematicians often seem to believe – just because they can prove hard theorems – that they are beyond the reach of fallacies of human nature, and often underestimate human psychology or think they are above it.

LET US IMAGINE A SCENARIO (which will not happen, since there is resistance: you are reading one voice in that resistance) in which all mathematicians drop human thinking and start using AI to write papers, from conception to publication. → This leads to an exponential production of manuscripts that no human could ever verify. → As a consequence, one of the main pillars of modern mathematics – to *trust* that a result is correct – is delegated to automated proof assistants. → As a consequence, another main pillar of modern mathematics – conveying *understanding* and not mere production – is destroyed: everyone quickly loses interest and motivation to produce any remaining human research (note *human psychology*). → The role of a current professional mathematician with a permanent position reduces to teaching. → With incentives for human-done mathematics disappearing, the number of new undergraduate students converges to zero. → All hiring of new academic staff in mathematics stops. → Human competence in mathematics disappears quickly. (This is already happening in some other fields where AI automation is the current dogma.) → In less than ten years there is not a single human who can tell you what the current state of the art in any branch of mathematics is.

THE ALTERNATIVE, which I believe is already happening, is the following. A split is occurring: between mathematicians who refuse working with AI – group $\mathcal{A}$ – and those who embrace it fully – group $\mathcal{B}$. (There is a third group of those who claim they take no sides, but that position will not be tenable for much longer.) $\rightarrow$ Say you are in group $\mathcal{B}$. Once you put your name to a paper which is partially AI-produced, your reputation (which is a large part of an identity of a mathematician) disappears: even if you subsequently publish a manuscript on which you had worked alone for years, no one will believe you produced it yourself (*human psychology* again). $\rightarrow$ It is impossible to switch back to group $\mathcal{A}$ (since your reputation as a thinking mathematician is gone), so you double down in the group $\mathcal{B}$. $\rightarrow$ Students are not stupid: they know which mathematician uses their brain and which, instead, uses a prompt: as a consequence, you lose students and postdocs who do not want to spend time playing with prompts. $\rightarrow$ As no one is interested in your AI-generated results, you lose interest in mathematics (*human psychology* again). $\rightarrow$ Group $\mathcal{A}$ ignores AI-generated mathematics and continues, to some extent, as before the advent of AI. $\rightarrow$ The population of professional mathematicians shrinks and an equilibrium settles: this is similar to an equilibrium discussed in [9], but I believe it will happen somewhat earlier than in that description.

If you use AI to generate papers. If you use AI to generate papers, I believe the following will happen.

IF YOU HAVE TENURE: Your reputation among mathematicians who do not use AI will disappear and will not be rebuilt again; and among mathematicians who use AI, no one cares about “your” results, since everyone is interested only in “their own” results. Your competence in mathematics will weaken; you will lose students and postdocs. You will “collaborate” only with other AI-prompts. The job will be psychologically unsatisfying and you will lose interest in mathematics completely.

IF YOU HAVE A PHD BUT NOT TENURE: All of the above applies. Additionally, you will never get tenure, unless you are being pushed by someone in your camp who has tenure. Regardless of how good you are in prompting a machine to produce you a paper, your mathematical skills will not develop out of passive learning. You will feel that you need to produce more output to compensate for lack of competence: but increasing the volume is counter-productive. People with tenure who encourage you to use AI in your research are not helping you get tenure. You should stop following their advice, but by now it is too late.

IF YOU ARE WORKING TOWARDS A PHD: All of the above applies. Additionally, you will never obtain a doctorate, unless your institution disregards good research practice and awards you a PhD for work that everyone knows is at least partly produced by a machine. People with a doctorate, who encourage you to use AI in your research, are not helping you get a PhD. You should stop following their advice, but by now it is too late.

IF YOU ARE A BACHELOR OR MASTER STUDENT: Read the above and derive a conclusion for your own mathematical path.

Why mathematics? Many propose the following explanation for why there is substantial “progress” in the use of AI in mathematics: because mathematics is a closed system, which does not rely on outside environment (such as the nature) to verify whether its statements are true; and that there are many employees of AI companies who can appreciate, to some extent, mathematical theorems and conjectures.

Whereas this is certainly true, the main reason is economics: AI companies have not even remotely delivered on promises they started from: cures for illnesses, climate change, nuclear fusion; instead there are large investments in powerful hacking tools and military use. In absence of success on problems they had advertised, they seem to believe that solutions to abstract mathematical problems is a suitable compensation: they have to show that they are *doing* something which seems meaningful, and they found a fruitful ground among a group of mathematicians. If you use AI in your research, you are giving AI companies free advertisement, a sense of purpose and a lifeline. Why?

Why? The question that I keep asking myself ever since the AI saga started is: Why? If you are a working mathematician, and you accept that the long-term impact on mathematics as a profession and on your own life is detrimental, and you still use AI to produce papers: Why?

Because everyone else you know is doing it? Because it is “comparable to the advent of search engines in the early 2000s”? Because it is “a new industrial revolution”? I suggest you actually read about the industrial revolution and learn how it ate its own children, for a century.

Epilogue: on farce, irony, shame. I finish with a selection of quotes.

In the text [10] that tried to “digest” an AI generated counterexample to a conjecture, one reads, beyond farce: *‘I used an AI chatbot to discuss various aspects of this problem and to confirm several of the calculations made here.’*

In [3] one reads: *‘... any PhD student who was not paying \$200 per month to access these tools was crazy’*, and ends, irony-oblivious: *‘What a time to be alive.’*

In [11] one reads: *‘When I was a postdoc, I went to the library once and it felt like a traumatic experience. I decided, “I’m never going back to the library.” There was no point.’* There were times when mouths would have been ashamed to say this, publications ashamed to print it, and eyes ashamed to see it. In our times one receives a medal instead.

REFERENCES

- [1] *Where I stand on AI*, <https://www.uni-saarland.de/lehrstuhl/lazic/prof-dr-vladimir-lazic/vladimir-lazic-where-i-stand-on-ai.html>.
- [2] *A recent experience with ChatGPT 5.5 Pro*, <https://gowers.wordpress.com/2026/05/08/a-recent-experience-with-chatgpt-5-5-pro/>.
- [3] *Human mathematicians are being outcounterexampled*, <https://xenaproject.wordpress.com/2026/07/20/human-mathematicians-are-being-outcounterexampled/>.
- [4] *How Terry Tao became an evangelist for AI in math*, <https://www.quantamagazine.org/how-terry-tao-became-an-evangelist-for-ai-in-math-20260608/>.
- [5] *Sometimes being first means seeing the end before anyone else*, <https://www.quantamagazine.org/jacob-tsimmerman-wins-2026-fields-medal-for-andre-oort-conjecture-proof-20260723/>.
- [6] *The last generation of mathematicians*, <https://www.youtube.com/watch?v=6uIJdXmB4vE>.
- [7] M. Weinreich, *The crisis of AI-generated mathematics*, <https://arxiv.org/pdf/2608.02859>.
- [8] *Code of Conduct “Safeguarding Good Research Practice”*, https://zenodo.org/records/14281892?preview_file=en.Guidelines+for+Safeguarding+Good+Research+Practice.Code+of+Conduct.Version+1.2.pdf.
- [9] ———, *Proof inflation*, <https://tjoresearchnotes.wordpress.com/2026/07/28/proof-inflation/>.
- [10] *A digestion of the Jacobian conjecture counterexample*, <https://terrytao.wordpress.com/2026/07/21/a-digestion-of-the-jacobian-conjecture-counterexample/>.
- [11] *Jacob Tsimerman on getting to the fun faster with AI — and worrying about the future*, <https://www.ams.org/journals/notices/202607/noti3372/noti3372.html>.