

Today, Levent Alpöge and I have made public three results: finite-time blowup with smooth forcing for incompressible porous media, for Boussinesq, and for 3d incompressible Euler.

We believe we also have blowup for hypo-dissipative Navier-Stokes. We are not releasing that paper today: unlike the above, the Lean verification has not yet finished. We do not yet have anything resembling a presentable writeup. I mention it because it is suggestive of a path to unforced Euler.

The program this fits into was not started by us nor was it proposed by a Large Language Model. The credit for the basic idea of this program goes to Diego Córdoba and Luis Martínez-Zoroa, who for several years have been exploring the construction of forced blow ups. We took their work as a starting point, using Large Language Models to push their program to completion.

Concretely, what Levent and I did was to take the Córdoba and Martínez-Zoroa program, which achieved blowup results with rough forcing, and, with a great deal of help from LLMs, push it to smooth forcing and to the incompressible Euler equations. The ideas making this line of attack possible are due to Córdoba and Martínez-Zoroa. Let me make plain what I have said to colleagues in private: in view of this body of work, I believe Luis Martínez-Zoroa deserves a Fields Medal.

My work with Levent has been a purely personal collaboration, free of any institutional agreements or official involvement by either of our employers. We used several LLMs throughout: Anthropic's Claude, OpenAI's Codex, especially with GPT-5.6 Sol and, more recently, Astra. The latter was only used for writeups and auditing our arguments.

For most of the past year progress was slow. We worked through the literature and upgraded various preliminary results, up to obtaining finite time blow up for the Incompressible Porous Media equation (with smooth forcing). This was until about a month ago, when we had real progress: on August 15th, we obtained the blow up results, with smooth forcing, for both Boussinesq and Euler. I can say the first LLM generated proof Levent sent me was the most horrendous I have ever read; we verified it on Lean on August 22nd. Since this point, we have been working around the clock to understand this proof and turn it into something readable.

There is another part of this story, and one that, honestly, I very much wish I did not have to be concerned with. I am not happy about the presentation quality in these papers. Ideally, we would have preferred to spend weeks turning the LLM generated proofs into something readable from the very first page. This level of care is what these problems and the community devoted to these problems deserves. The various Boussinesq and Euler write-ups in particular are much closer to what models produce under human direction than to a paper written by a person. The Euler writeup, in particular, can only be described as AI slop. I am sorry for this. The reasons are below, and they involve our being pressured by outside factors. I say this not as an excuse but as an explanation.

I had planned to say on announcing our work that the results are not the important thing. Rather the important thing is instead the significance that a mathematician and an LLM model can now do all this work in a month. The

significance of this with respect to the way we train students, assign credit, referee, and decide what is worth one human life's attention cannot be understated. This is a Deep Blue-Kasparov moment. The community needs to have serious and unhurried discussion about where to go from here. Instead of these incredibly important developments, I find myself writing about something else. I want to set out what happened as plainly as I can.

On Thursday, September 3rd, with a rumor circulating that Anthropic had resolved a major open problem, and with Levent having received tips that information about our progress had been passed to OpenAI, I wrote to a prominent mathematician at OpenAI. I am quoting my email in full because I would rather the full text be read rather than my summary of it:

"Hi [...],

We have not met in person, but I was one of the speakers at the [...].

I am writing because a rumor seems to be spreading quickly that Anthropic has resolved a major open problem. I cannot be certain I am the person it attaches to, but a colleague at Courant emailed me about it yesterday – who heard it from an analyst in the UK, who had it from somewhere further upstream – so it seems safe to assume I am. I gather versions linking Levent Alpöge to a solution are going around in tech as well. There however appears to be a lot of confusion with regards to the exact problem solved.

I should also emphasize that this is not an institutional effort. It is a strictly personal collaboration between the two of us, and there is no formal agreement behind it. I pay for the tools my group uses out of my own research funds, including footing a large bill to OpenAI. I have had an industry collaboration before, with DeepMind, which had a formal institutional arrangement.

We do have work in this area that we are confident in, and we will post it shortly, the paper and the formalization together. We intentionally decided against rushing out a Lean certificate alongside an unpolished preprint. I feel strongly that the first thing anyone reads should be a mathematical argument presented in the normal manner, rather than just a formal certificate.

I am writing to you directly rather than saying anything publicly so that you have the facts to address this on your end.

Best wishes,
Tristan"

He replied the same day:

"If you are willing to give any details it would be useful to avoid competing here and in general we are always thrilled when mathematician make progress with our models. Additionally if there is anything in terms of compute from OpenAI's end we would be happy to provide it."

I asked to speak the following week. On Friday, September 4th, I was asked whether I could meet that day; I again said the following week. At 12:45 on Sunday, September 6th, I was asked whether I could meet "at any point today." Sebastien Bubeck joined. The three of us spoke twice that afternoon. Levent was not on the calls.

I was told that an internal OpenAI model had produced a proof of finite time blowup for the forced Navier-Stokes equations. When Levent asked by text for

the precise statement, the answer was: “Existence of forced blowup in R^3 and T^3 ,” with “the forcing function is smooth option c and d in Fefferman.” I was told the proof is about 100 pages. I have not seen it.

I should say here why I interpreted their statement the way I did, the interpretation I will discuss below. The route to the Clay problem through a smooth force, options c and d in Fefferman’s statement of the problem, is the route Luis and Diego opened and the one Levent and I had quietly chosen to attack. Almost nobody else I know of was working on it. It is not the direction one arrives at in a few days by giving a model the problem statement. When I heard “forced,” it was a bright red flag.

I was shown a prompt and told the internal research model had simply been given the problem statement. Levent had been told by Sebastien “very little human input” had been used. This turned out not to be true. Over the course of the call, as members of their team sent Sebastien corrections and details over their internal chat, it emerged that an entire team had been working on the problem, that this was one of a number of things that was tried, that work had started on the unforced problem, that the team first set the model on easier problems, including Euler, that even the prompt that had been shown to me had been written by prompting Codex, and that an insane amount of compute had been used.

I asked when the first prompt had been sent by them. This question was not answered directly by OpenAI for some time. Eventually it was agreed that it had been sent in the past few days, after information about our work had reached OpenAI.

I asked whether the model had been trained on, or had access to, our sessions in Codex, into which we had been putting all our drafts for the whole of this project. I was told the model did not look up user data. I asked again, about training, and I did not get an answer.

Two proposals were offered to me. The first was that we post our Euler result, and that OpenAI post its Navier-Stokes result the next day. The second was that, after posting Euler, I alone write a paper presenting the Navier-Stokes result, acknowledging that an internal OpenAI model had resolved it. Sebastien twice asserted that he wanted Levent removed from authorship, and said it would all be simple if only it were not the case that, and it was so annoying that, Levent works at Anthropic. It was also said that if OpenAI posted after us, they would say that we deserved the Clay Prize, and that we were the “closest humans to the problem”. I declined both offers.

I said that if OpenAI released its result in the way proposed I would go public with what happened. The reply was, “Why would you ruin your career?” I replied that I am an academic, and asked why he thought going public would ruin my career. The reply was, “If you don’t want me to be nice, then I don’t have to be nice.”

Some time later Levent received a text proposing that he and Sebastien speak one on one, saying, “I don’t know if Tristan is being fully rational right now.” Levent declined and said conversations should be with me. Sebastien sent a follow-up email that night requesting to speak with me on Monday, September

7th. I did not respond.

I would like to be clear about what I am not claiming. I have not seen OpenAI's proof. I do not know what their model did, or how. I do not know whether our data was used. I am not accusing anyone of anything. I am stating what I was told, when, and what was proposed to me. I am stating it because the alternative is to let a sequence of announcements say something I know to be false.

If indeed an OpenAI model did close the gap to Navier-Stokes, that is a remarkable thing and it should be said loudly, by them, with the history intact.

I would much rather be talking about mathematics, Luis and Diego's ideas, and what this all means for the rest of us.

Lastly, I would like to thank the entire mathematics community that have been so supportive of me over the last 24 hours.