

GUEST ESSAY

The Fever Dream of Imminent Superintelligence Is Finally Breaking

Sept. 3, 2025

Listen to this article · 8:26 min [Learn more](#)

By Gary Marcus

Mr. Marcus is a founder of two A.I. companies and the author of six books on natural and artificial intelligence.

GPT-5, OpenAI's latest artificial intelligence system, was supposed to be a game changer, the culmination of billions of dollars of investment and nearly three years of work. Sam Altman, the company's chief executive, implied that GPT-5 could be tantamount to artificial general intelligence, or A.G.I. — A.I. that is as smart and as flexible as any human expert.

Instead, as I have written, the model fell short. Within hours of its release, critics found all kinds of baffling errors: It failed some simple math questions, couldn't count reliably and sometimes provided absurd answers to old riddles. Like its predecessors, the A.I. model still hallucinates (though at a lower rate) and is plagued by questions around its reliability.

Although some people have been impressed, few saw it as a quantum leap, and nobody believed it was A.G.I. Many users asked for the old model back.

GPT-5 is a step forward but nowhere near the A.I. revolution many had expected. That is bad news for the companies and investors who placed substantial bets on the technology. And it demands a rethink of government policies and investments that were built on wildly overinflated expectations. The current strategy of merely making A.I. bigger is deeply flawed — scientifically, economically and politically. Many things, from regulation to research strategy, must be rethought. One of the keys to this may be training and developing A.I. in ways inspired by the cognitive sciences.

Fundamentally, people like Mr. Altman, the Anthropic chief executive Dario Amodei and countless other tech leaders and investors had put far too much faith into a speculative and unproven hypothesis called scaling: the idea that training A.I. models on ever more data and using ever more hardware would eventually lead to A.G.I. or even a superintelligence that surpasses humans.

However, as I warned in a 2022 essay, “Deep Learning Is Hitting a Wall,” so-called scaling laws aren’t physical laws of the universe like gravity but hypotheses based on historical trends. Large language models, which power systems like GPT-5, are nothing more than souped-up statistical regurgitation machines, so they will continue to stumble into problems around truth, hallucinations and reasoning. Scaling would not bring us to the holy grail of A.G.I.

Many in the tech industry were hostile to my predictions. Mr. Altman ridiculed me as a “mediocre deep learning skeptic” and last year claimed “there is no wall.” Elon Musk shared a meme lampooning my essay.

It now seems I was right. Adding more data to large language models, which are trained to produce text by learning from vast databases of human text, helps them improve only to a degree. Even significantly scaled, they still don’t fully understand the concepts they are exposed to — which is why they sometimes botch answers or generate ridiculously incorrect drawings.

Scaling worked for a while; previous generations of GPT models made impressive advancements compared with their predecessors. But luck started to run out over the past year. Mr. Musk's A.I. system, Grok 4, released in July, had 100 times as much training as Grok 2 had, but it was only moderately better. Meta's jumbo Llama 4 model, much larger than its predecessor, was mostly also viewed as a failure. As many now see, GPT-5 shows decisively that scaling has lost steam.

The chances of A.G.I.'s arrival by 2027 now seem remote. The government has let A.I. companies lead a charmed life, with almost zero regulation. It now ought to enact legislation that addresses costs and harms unfairly offloaded onto the public — from misinformation to deepfakes, A.I. slop content, cybercrime, copyright infringement, mental health and energy usage.

Moreover, governments and investors should strongly support research investments outside of scaling. The cognitive sciences (including psychology, child development, philosophy of mind and linguistics) teach us that intelligence is about more than mere statistical mimicry and suggest three promising ideas for developing A.I. that is reliable enough to be trustworthy, with a much richer intelligence.

First, humans are constantly building and maintaining internal models of the world — or world models — of the people and objects around them and how things work. For example, when you read a novel, you develop a kind of mental database for who each character is and what he or she represents. This might include characters' occupations, their relationships to one another, their motivations and goals and so on. In a fantasy or science fiction novel, a world model might even include new physical laws.

Many of generative A.I.'s shortcomings can be traced back to failures to extract proper world models from their training data. This explains why the latest large language models, for example, are unable to fully grasp how chess works. As a result, they have a tendency to make illegal moves, no matter how many games they've been trained on. We need systems that don't just mimic human language; we need systems that understand the world so that they can reason about it in a deeper way. Focusing on how to build a new generation of A.I. systems centered on world models should be a central focus of research. Google DeepMind and Fei-Fei Li's World Labs are taking steps in this direction.

Second, the field of machine learning (which has powered large language models) likes to task A.I. systems to learn absolutely everything from scratch by scraping data from the internet, with nothing built in. But as cognitive scientists like Steven Pinker, Elizabeth Spelke and me have emphasized, the human mind is born with some core knowledge of the world that sets us up to grasp more complex concepts. Building in basic concepts like time, space and causality might allow systems to better organize the data they encounter into richer starting points — potentially leading to richer outcomes. (Verses AI's work on physical and perceptual understanding in video games is one step in this direction.)

Finally, the current paradigm takes a kind of one-size-fits-all approach by relying on a single cognitive mechanism — the large language model — to solve everything. But we know the human mind uses many different tools for many different kinds of problems. For example, the renowned psychologist Daniel Kahneman suggested humans utilize one system of thought — which is quick, reflexive and automatic and driven largely by the statistics of experience but is superficial and prone to blunders — along with a second system that is driven more by abstract reasoning and deliberative thinking but is slow and laborious. Large language models, which are a bit like the first system, try to do everything with a single statistical approach but wind up unreliable as a result.

We need a new approach, closer to what Mr. Kahneman described. This may come in the form of neurosymbolic A.I., which bridges statistically driven neural networks (from which large language models are drawn) and some older ideas from symbolic A.I. Symbolic A.I. is more abstract and deliberative by nature; it processes information by taking cues from logic, algebra and computer programming. I have long advocated a marriage of these two traditions. Increasingly, we are seeing companies like Amazon and Google DeepMind take such a hybrid approach. (Even OpenAI appears to be doing some of this, quietly.) By the end of the decade, neurosymbolic A.I. may well eclipse pure scaling.

Large language models have had their uses, especially for coding, writing and brainstorming, in which humans are still directly involved. But no matter how large we have made them, they have never been worthy of our trust. To build A.I. that we can genuinely trust and to have a shot at A.G.I., we must move on from the trappings of scaling. We need new ideas. A return to the cognitive sciences might well be the next logical stage in the journey.

Source images via Getty Images.

Gary Marcus is a professor emeritus at New York University and was a founder and chief executive of Geometric Intelligence. His most recent book is “Taming Silicon Valley.” He publishes a newsletter about A.I.

The Times is committed to publishing a diversity of letters to the editor. We’d like to hear what you think about this or any of our articles. Here are some tips. And here’s our email: letters@nytimes.com.

Follow the New York Times Opinion section on Facebook, Instagram, TikTok, Bluesky, WhatsApp and Threads.