

Print Edition

Saturday 1 August 2026

EDITION NO. 9 • 2 PIECES • 9 PAGES

IN THIS EDITION

- | | | |
|----|---|---|
| 1. | What will more intelligence actually do for us? • Noahpinion | 1 |
| 2. | Why we're epistemically flexible but stubbornly self-regarding • The Update | 8 |
-

What will more intelligence actually do for us?

BY NOAH SMITH • NOAHPINION • 26 JULY 2026



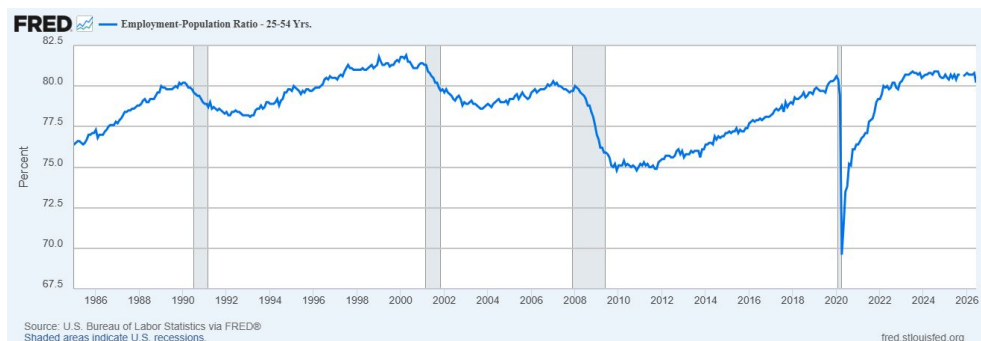
“[I]ntelligence, that counterentropic conjoined twin of information, must become the most powerful force in the universe, the energy to which all other physical laws must eventually kneel...Intelligence was destiny, manifest.” — Ian McDonald, “Verthandi’s Ring”

Not a lot of people expected that AI would come for the mathematicians before it came for the truck drivers, but it did. The other day, an AI model disproved the Jacobian Conjecture — an 87-year-old open problem that human mathematicians had struggled to solve. The greatest living human mathematician, Terence Tao, turned to AI to help him understand the solution. Around the same time, AI solved a very important open question in quantum cryptography. Solving Erdos problems has now become almost child’s play for the best AI models. And this is the worst AI will ever be at *math*. Model capabilities, and the amount of compute available, both

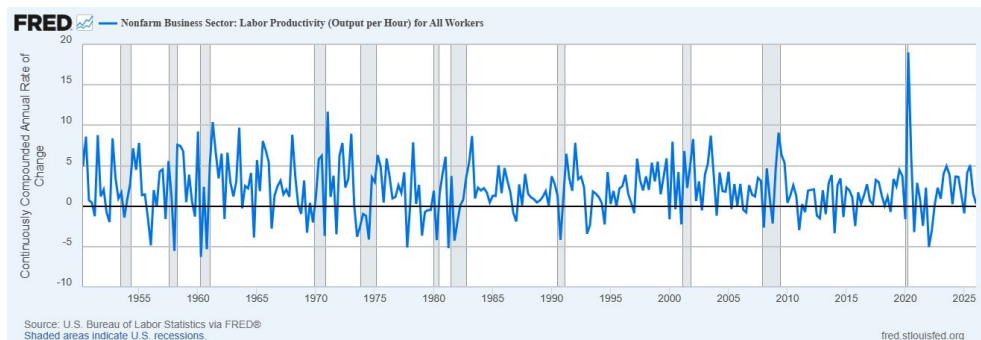
continue to increase at rapid rates. (Meanwhile, long-distance trucking employment is slightly higher than it was a decade ago.)

I don't expect mathematicians to actually lose their jobs en masse, of course.¹ But it's becoming clearer and clearer that humankind has invented machines that are *smarter* than we are. Intelligence isn't defined for machines the same way it is for humans — AI's capabilities are spiky in different ways than ours — but it's undeniable that the technology is improving rapidly in every domain of cognitive capability. It's still possible to find some mental tasks that humans are better than machines at, but those final advantages tend to disappear almost as quickly as we can identify them. "AGI", or "ASI", or whatever you want to call it, is certainly here.

And yet...the world remains much the same. In lots of sci-fi books, as soon as artificial superintelligence arrives, it bootstraps itself to even more godlike intelligence in an explosive "singularity" that rapidly transforms the entire physical universe. Lots of people, especially "AI safety" and "effective altruist" types, expected things to play out basically the same way in reality. But looking around, not much has changed since we entered the intelligence explosion. There's a huge data center boom, and most people use AI on a daily basis, but we still live basically the same lives — driving to work or taking the train, sitting in front of a computer, scrolling on our phones, collecting a paycheck. People are staying in their jobs longer, but employment hasn't been disrupted in a significant way:



Meanwhile, we've had decently robust productivity growth, but nothing really amazing:



A lot of people I know are surprised by this. Ruxandra Teslo writes:

Walking around the world today one might notice that it is weirdly unchanged...To many, this is surprising. Just the other day I was at a conference where someone remarked that if he could have seen today's AI capabilities a few years ago, he would have been astonished — and would have assumed the world by now would look far more transformed, with much higher GDP growth.

And Clifford Sosin writes:²

Superintelligence arrived. You probably didn't notice, because it turned out to be kind of incremental...Don't believe me? Run the test. Talk to Fable 5 for an hour, then talk to your ten smartest friends. Which one is smarter? Don't worry, they won't be offended...We were told to expect something bigger. Once machines crossed some line, the system's IQ would climb to heights we couldn't follow, and we'd be sharing the planet with something that designs warp drives and thinks thoughts as far past us as mine are past my dog...What we got is a tool that writes excellent code, beats people at a startling range of tasks, and will clearly reshape the economy. It's also, somehow, incremental. No takeoff. No explosion. That's the strange part.

Teslo blames bottlenecks — governance and other “frictions” — for the slow economic impact. But some others are advancing a more radical hypothesis³ — that intelligence itself is subject to diminishing returns.

One of these is Francois Chollet, an AI researcher who specializes in measuring AI's capabilities. In a highly controversial series of tweets back in March, he conjectured that intelligence might be subject to diminishing returns:

One of the biggest misconceptions people have about intelligence is seeing it as some kind of unbounded scalar stat, like height. "Future AI will have 10,000 IQ", that sort of thing. Intelligence is a conversion ratio, with an optimality bound. Increasing intelligence is not so much like "making the tower taller", it's more like "making the ball rounder". At some point it's already pretty damn spherical and any improvement is marginal.

Now of course smart humans aren't quite at the optimal bound yet on an individual level, and machines will have many advantages besides intelligence -- mostly the removal of biological bottlenecks: greater processing speed, unlimited working memory, unlimited memory with perfect recall... but these are mostly things humans can also access through externalized cognitive tools.

In fact, this is a possibility I myself had raised in a post a year earlier:

It seems possible that humans are simply incredibly specialized in a few types of cognitive tasks — extracting patterns from sparse data, synthesizing various patterns into “intuition” and “judgement”, and communicating those patterns in language — and that we've basically approached the theoretical maximum in those narrow

areas...That would explain why AI has gotten much better at things like math and coding and forecasting over the last year, but why the basic chatbot interface doesn't seem much more "intelligent". It would also explain why when you talk to Terence Tao about math, it's like talking to a superhuman, but when you talk to him about where to get lunch or which movies are the best, he'll just sound like a fairly smart normal dude. AI will eventually get better than Tao at math...but it may never get much better than the most thoughtful, eloquent humans at deciding where to get lunch or recommending movies. It may simply not be mathematically possible to get much better than we already are at that sort of thing.

Why would intelligence top out like this? Well, if we think of intelligence as *the ability to extract information from data*, then even an infinitely advanced model endowed with infinite compute will be limited by the fact that *there's a limited amount of information that can be extracted from the data*.

For one thing, data itself is in limited supply. You can't transform the world unless you can (in some generalized sense) *understand* it, and you can't understand the world unless you can *measure* it, and our ability to measure the world is inherently limited and finite.

This is basically the hypothesis advanced by Arvind Narayanan and Sayash Kapoor:

We think there are relatively few real-world cognitive tasks in which human limitations are so telling that AI is able to blow past human performance (as AI does in chess). In many other areas, including some that are associated with prominent hopes and fears about AI performance, we think there is a high "irreducible error"—unavoidable error due to the inherent stochasticity of the phenomenon—and human performance is essentially near that limit....We predict that AI will not be able to meaningfully outperform trained humans (particularly teams of humans and especially if augmented with simple automated tools) at forecasting geopolitical events (say elections).[emphasis mine]

Sosin says something similar:

We hold a thin scatter of facts about the world, and intelligence or reasoning is whatever fills the space between them...Where the space between the facts behaves well, this is close to godlike...Coding, math and most administrative work are [like this]. What makes them easy is that they have relatively smooth solution spaces and are tractably verifiable...Most of what matters doesn't behave like that. The universe is mostly the emergent behavior of complex systems...The limit is contact with reality. A smarter reasoner fills the gaps between known facts in simple areas faster and better, but it doesn't produce new facts.

Sosin makes an important point here, which is that the limitations of intelligence aren't necessarily about limited data. Even if we can keep on

collecting infinite data, the cost of extracting additional information from that data might explode to infinity. This is the idea of *chaos*. Even in a deterministic universe where the present states of all particles are enough to perfectly determine the future, our ability to predict the future can be inherently limited; the tiniest infinitesimal error in our measurement of the present explodes into a huge error when we attempt to extrapolate even a small distance into the future.

So although we don't know yet, it's possible that humans were already hitting the point of diminishing returns with regards to individual cognitive capacity, and that superintelligent machines will never be as far beyond us as we are beyond dogs. But even if that's true, I can think of at least three reasons why machine superintelligence could still deliver huge productivity gains.

Smart matter

The most obvious advantage that machine superintelligence confers is replicability. The number of human intelligences is fixed by the fertility rate, and we don't know how to substantially boost that rate; in fact, it's falling inexorably, and the human race is set to shrink.

AI isn't bound by those limitations. By building more data centers with more compute, you can run more agents in parallel — essentially, you get more cognitive work. It's not free, but nor is it limited. It's good old physical capital; unlike human capital, you can just build more of it whenever you like, simply by reinvesting some portion of your economic output. Imagine if we suddenly discovered a way to manufacture more *land* in any city on the planet; this is similar.⁴

Of course, lots of tasks are physical ones, and for these you need physical machines — basically, robots. This is probably why so many AI people are now working on robotics, world models, physical AI, and so on. The roboticization of the world has a huge tailwind — the battery revolution, which allows energy to be stored and moved around much more easily. Conveniently, we got the physical tools to turn dumb matter into smart matter just as we also got the digital tools.

(It's also worth noting that like the energy in batteries, the intelligence in robots is divisible. A robot the size of an ant can be remotely controlled by a data center the size of a football field. This is also a capability that human intelligence lacks.)

This doesn't mean economic output will explode to infinity. But what it does mean is that humans will be able to use physical capital — GPUs and robots — to do more and more tasks at once, including many cognitive tasks that we used to do the hard way. In the long-run steady state, this should increase the capital-to-labor ratio of our society; each human will essentially leverage an army of intelligent machines. It's basically another industrial revolution, and it has very little to do with whether artificial intelligence is *smarter* than human intelligence in any sort of head-to-head matchup.

Distributed tacit knowledge

The German company Zeiss makes the best glass on the planet. If one of the mirrors that Zeiss makes for ASML's EUV chipmaking machines were the size of Germany, the biggest bump on that mirror would be just one millimeter high. Only a few other companies — and maybe no other company on Earth — can match that. Zeiss' mirrors also have a number of other amazing properties, like not distorting much due to temperature changes.

How does Zeiss make glass this good? No one knows — not even the people at Zeiss. If the technology were capable of being written down on a blueprint, China would have hacked Zeiss and stolen it, the way Huawei hacked Cisco and Nortel. If the technology were capable of being explained by a former Zeiss employee, or even several former Zeiss employees, China would have paid those people many millions of dollars to spill the beans.

Zeiss' technology basically can't be stolen, because it's tacit and distributed. It consists of a vast number of little tricks and techniques that a huge number of individual employees use on a daily basis. These people don't always even realize all those little things they're doing that make the glass come out so good. And each employee knows a different set of tricks and techniques. The knowledge exists at the level of the organization itself, and is thus very hard to steal or recreate.

This is true of lots of corporate technology. A big part of the reason China can cut off the supply of rare earths to the rest of the world any time it wants to is that other countries aren't very good at refining rare earths. Rare earths are difficult to separate from each other in solutions; it takes a ton of little chemistry tricks to do it cheaply at scale. Chinese refiners have spent four decades building up those little tricks and techniques; American or Japanese refiners won't simply be able to replicate their efficiency overnight, and so it'll continue to cost much more to produce rare earths outside China.

Except in the age of AI, this might change. Suppose American rare earth refiners give their employees a bunch of equipment to record everything they do — smart glasses, gloves, and so on — in addition to sensors distributed throughout their plants. AI will be able to synthesize all that information and very rapidly suggest small ways to improve the production process. Many of those little experiments will fail; others will succeed and will quickly be adopted, allowing another round of experimentation and improvement to begin very quickly. Crucially, AI's ability to do this doesn't depend on its raw intelligence — only on its ability to handle huge amounts of data very quickly.

In other words, in the age of AI, distributed tacit knowledge might not be nearly as big of a barrier to technological diffusion. This could improve economy-wide productivity, as lagging firms catch up to leading firms much more quickly. A more equal distribution of productivity would also make the economy more competitive, creating more surplus for consumers (though possibly reducing the incentive for firms to innovate, by making technology less excludable).

AI's ability to quickly produce distributed tacit process knowledge might also supercharge productivity growth at the frontier. Imagine if any company could optimize any production process five times faster than today. The whole economy would speed up, as components got cheaper, turnaround times and product cycles got shorter, and scale-up got much faster.

And as with the previous example, improving the production of distributed tacit knowledge wouldn't depend on AI's raw intelligence. It would spring from AI's ability to act like a computer — to interface directly with sensors, to handle lots of data, to perceive tiny details, and to do everything very very quickly.

Cloud laws

For decades, researchers in the field of natural language processing tried to figure out the principles behind human linguistic communication. They made frustratingly little progress; the processes by which humans convey information to each other through words just don't seem to obey simple laws, like the ones that govern electromagnetism or the circulatory system.

Then along came AI, and suddenly linguistic communication seemed like a solved problem. LLMs can reliably sound like a human being, even if we don't understand how they manage to do it.

What if there are lots of other aspects of the Universe that work the same way — too complex to understand in terms of simple laws, but not so complex that they just dissolve into unknowable chaos? It's possible that we can reliably control these complex phenomena with AI, even if we never reduce them to the kind of principles that we can teach a grad student in a textbook. In fact, I wrote an essay called "The Third Magic", where I suggested that this might be equivalent to a whole new scientific revolution:

Another way of saying this is that there may be laws of the universe that humans can't understand but AI can. I call these "cloud laws" — causal regularities that can be exploited by technology, but which are too diffuse and complex for an individual human being to either intuit or communicate.

Human language seems to obey cloud laws, so why not other phenomena too? Perhaps social sciences like economics, sociology, and political science obey similarly complex regularities, and AI can help us find them. Perhaps there are physical processes — plasma, or topological materials, or aerial turbulence, etc. — that obey cloud laws instead of chaos?

In other words, thanks to AI, we might be on the precipice of a new age of scientific advancement. And this won't necessarily depend on how smart AI is in comparison to a single human; it'll depend on its computer-like ability to hold huge amounts of data in its working memory and extract complex patterns from that data.

If this turns out to be true, it means Francois Chollet is wrong. Chollet hypothesizes that groups of humans, using pre-AI computing tools, can approach AI's level of scientific competence. But human collaboration is bottlenecked — it's limited by our ability to intuit patterns at an individual level, and to communicate these patterns from one individual to another.

er. AI, being a computer, just doesn't have this sort of limitation; it can work with vast, diffuse patterns without having to break them into pieces or simplify them in order to compress them into the tiny pipelines of person-to-person explanation.

So even if AI never gets much better than humans at the kind of science that humans have done heretofore, it might open up whole realms of scientific discovery that have previously been totally inaccessible to even the largest groups of the smartest humans. If much of the Universe turns out to be ruled by cloud laws, we could be on the precipice of a scientific renaissance.

The common thread in all three of these examples is that AI may revolutionize productivity not by being much smarter than a single individual human — not by simply solving harder and harder math problems — but by marrying human-style intelligence to the vast, inhuman capabilities of computers. We could simply be thinking about the benefits of intelligence wrong — arrogantly privileging the kind of mental tasks we humans happen to do especially well, while ignoring the value of the tasks we do poorly.

1

Why? Several reasons. Tenure still exists. Humans who can do math at a high level will still be needed if we care about *understanding* the results that AI spits out. Human math teachers will still probably be valuable. And most importantly, humans will be needed to tell AI *what kind of math* we want it to solve, and why.

2

Or at least, prompts AI to write.

3

I think there is a widespread, tacit assumption among many intelligent humans that the quality that had made *them* stand out among their peers was the fundamental stuff of the Universe, the font of all value. I will write more about this at some point, but I think it helps explain why the idea that intelligence is just one production factor among many seems so unthinkable, heretical, and revolutionary to so many people in the tech industry.

4

Yes, I know we can reclaim land from the ocean. But the opportunities for this are fairly limited, and the cost is often very high.

Why we're epistemically flexible but stubbornly self-regarding

BY STEFAN SCHUBERT • THE UPDATE • 31 JULY 2026

Dan Williams recently published a great essay on two types of pessimism about humanity: moral or *motivational* pessimism (especially about overcoming selfishness) and *epistemic* pessimism – pessimism about our ability to acquire and deploy knowledge. He argues that Thomas Sowell unhelpfully bundles them, and that the case for motivational pessimism is considerably stronger (with some nuances).

I share Dan's views and highly recommend reading his essay. It got me thinking about why the motivational and the epistemic cases come apart. In my view, a root cause is that our inborn motivations are far more fixed than our inborn epistemic intuitions.

We can, of course, change our motivations to some degree, but there are clear limits. It's effectively impossible for people to become completely impartial between themselves and unrelated others. My sense is that even in communes organised around sharing, people remain intuitively selfish in much of their everyday behaviour.

By contrast, our epistemic intuitions are more flexible. Take, for instance, the belief that objects like stones, trees, and tables are solid – a view that science has, in a sense, overturned. We can easily accept that verdict, even though those objects will continue to appear solid to us. We can replace our intuitive worldview with one that's grounded in scientific evidence. In fact, we stake our lives on this revised worldview every time we board a plane or take a prescription drug.

Of course, many beliefs are distorted by cognitive biases and motivated reasoning, but their resistance to change is often overstated. Our beliefs form an interconnected web, so biased or self-serving beliefs will typically conflict with other things we believe. Moreover, we argue with others, who can reveal inconsistencies we refuse to see. This means that distorted beliefs can be corrected, and frequently are.

And we don't just change our beliefs, but also our methods for acquiring them. While we're naturally attracted to striking anecdotes, we've learnt that experiments and statistical inference are more reliable.

Even people with limited abilities can form fairly accurate beliefs thanks to a shortcut: deference to epistemic superiors. Instead of trying to figure things out themselves, they can listen to those with expertise. There's no similarly effective shortcut for improving the motivations of unusually selfish people.

The asymmetry between the motivational and the epistemic cases ultimately has evolutionary roots. Since the genes of unselfish individuals were less likely to spread, there was strong pressure to lock selfish motivations in place (as Dan notes). There was no comparable pressure to lock in specific beliefs, such as the belief that tables are solid. On the contrary, there was pressure towards a degree of epistemic flexibility, since so much of what we need to know can only be learnt from experience.