

MIT Just Revealed the AI Bubble's Fatal Flaw

22 MIN · YOUTUBE · [HTTPS://YOUTU.BE/3ESCLFR8M7I?IS=56PMSLIOWNGDZRIC](https://youtu.be/3ESclFR8m7I?is=56PMSLIOWNGDZRIC)

<https://youtu.be/3ESclFR8m7I?is=56PMSLIOWNGDzRic>

SUMMARY

The discussion centers on the intersection of wealth inequality and the artificial intelligence (AI) sector, particularly focusing on the impending IPOs of OpenAI and Anthropic. The speaker argues that the current hype surrounding AI is built on shaky foundations, and that the technology may not deliver the limitless potential that investors are betting on, drawing parallels to past market bubbles like Enron.

- *Wealth inequality is worsening, with AI at the center of this issue.*
- *The rush to IPO by OpenAI and Anthropic raises questions about the actual value of their technology.*
- *Historical parallels are drawn to Enron, where perceived limitless potential masked underlying financial issues.*
- *The fundamental question for AI companies is whether increasing model size will lead to better performance or if it is merely a speculative prophecy.*
- *Research indicates that open-source models are closing the performance gap with proprietary ones at a fraction of the cost.*
- *The speaker emphasizes the importance of critical thinking and understanding the real value behind AI technologies.*
- *If AI models do not improve significantly, the inflated valuations of AI companies could lead to a market correction.*
- *The narrative of "this time it's different" is a common theme in technology bubbles, often leading to eventual disillusionment.*

It's critical that we talk about wealth inequality and AI because the mania fueling the upcoming IPOs of OpenAI and Anthropic are based on what seems to be a lie. And it's one that we can disprove if we do one thing that most won't, which is bother to look. We are experiencing one of the largest divides in wealth and opportunity in the whole of human history. It's not getting any better. In fact, it's getting worse. >> Yeah. And in fact, it's no longer a straight line. K that the the arm of the K looks almost parabolic.

>> And at the heart of this issue is one technology artificial intelligence. >> There is an overinvestment in this abstract technology that we crave when confidence is extremely high. We imagine this unlimited opportunity in futuristic technology at every peak in confidence. And then there comes the backlash. >> And right now the two largest artificial intelligence companies in the world, OpenAI and Anthropic, are sprinting towards IPOs. So the question becomes, if this technology is so limitless that it simply cannot be valued in any normal terms, why rush to sell now? The reason is that there is a prophecy being sold, but behind the scenes, there's a fuse burning down. And a study from MIT reveals exactly how that bomb bursts.

And if we look at the implications, what we see is that this prophecy impacts our investments, our retirements,

our career decisions, and even how we educate our kids. And if we can't see the truth despite the prophecy, we make decisions in our lives that we can't undo. I'm Brendan Dell. This is the leverage class. Let's see through it. To understand what's happening with these two IPOs, we need to go back to the year 2000 and to this man, Jim Chanos, who has across his career done the thing that most won't. He looks so Chanos is a short seller, which said simply is someone who profits when a company stock goes down. Now, this makes him one of the most hated and the most useful people in finance because while everyone else gets paid when these prophecies keep going up, he only gets paid when he finds fundamental holes in the narrative. So, in 2000, he placed the most famous bet of his career, and it's one that directly parallels what's happening with AI right now. So, in 2000, Enron stock had more than quadrupled in value in 2 years. It was over 300% and it peaked around \$90 which made it the seventh largest company in America trading at more than 70 times earnings. So this was an energy company that Wall Street had decided was actually a technology company. The same costume that we work wore. We talked about this in a previous video. They said that Enron was too innovative, too complex, and too far ahead to value in any normal way, which would be like some reasonable multiple of what they actually sold or earned. But Chanos didn't buy that story. It started with a newspaper article. A friend pointed him to a piece about how these energy companies were booking their profits. So Chanos pulled Enron's actual financial statements and he found the hole. It was called gain on sale accounting and said plainly they counted money that they might earn as money that they had earned. It's like if you asked a bank for a loan based on a raise that you're sure that you're going to get next year.

So Chanos did something very simple. The math. For every dollar this company puts to work, how much does it actually earn back? The answer was seven cents. A 7% return before taxes. And it cost Enron somewhere between 7 and 9 cents to borrow each of those dollars in the first place. So the most admired company in America was making no money. In November of 2000, Channel started betting against it. And at that time, everybody basically laughed at him. An Enron executive even publicly thanked the short sellers, saying that they were keeping the stock cheap enough for employees to load up on more. Now, we all know how this ends. By December of 2001, Enron filed the largest bankruptcy in American history. At that point, the stock went to pennies. Chanos made around 500 million. Now, 25 years later, he is looking at AI and calling it worse than anything he's ever seen before. the diamond or platinum level of fraud, he calls it. And he thinks that this cycle is even riskier than the dot bubble because at least back then the companies buying the equipment were profitable like GE or AT&T. But today, all the companies driving demand are just lighting money on fire. Now, for their part, OpenAI and Anthropic would tell you the exact opposite of this story.

They would tell you this is the single most important technology in human history. AI is a concept. We can't define its parameters. And um it's a great concept. It's likely to be the most powerful force any of us have ever seen. But that's I think that's all we know. a decision to participate or not participate and what price to participate at is it's really uh well it's what my South African friends call a thumb said plainly the belief is that this technology is so powerful and infinitely scalable that there is no possible way to understand its value in the future. But what this perspective is missing is that there is one fundamental moat for these two companies which dictates whether or not they can live up to their hype and cover the debts being created and it boils down to one single word scale. So before we get to that a lot of folks out there are trying to decide what to do next in their careers.

I personally want three things out of work. I want work I enjoy doing. I want relative time freedom and I want diversified income far in excess of what I need. The way that I originally built this for myself was with consulting. And after many years, my clients were signing my S. So wondering how I've been able to put that business together. So I built some modules for them which started getting shared which I turned into a course called the freelance formula. It's a program for mid-career professionals who want to build their own independent business. So right now you can get that full program for \$99.

The link is below. With that, back to the content. So if you strip away all the noise around artificial intelligence in its present state, you are actually asking a much more direct perhaps not simple but direct question. Will building larger and larger and larger large language models continue to increase the performance of these models in a way that becomes one actual intelligence, two self-governing meaning that they become agentic and they can execute complex tasks on their own and three impossible to replicate without scale. That is the fundamental question of value because everything else the share of the enterprise or the MCP ecosystem or the third party builders the consumer adoption all those things that will ultimately drive the differentiated business value of the long term it hinges on that single question which brings up the crux of this entire movement. Can we understand the answer to that? Because if the answer is no, if bigger doesn't keep meaning better, then the price is in no way justified by what these companies are today. It's entirely betting on what they might become. And there is an old Wall Street line for pricing a company on a prophecy. Those who live by the crystal ball are destined to eat broken glass. And if that's what this is, a prophecy, the question stops being if and starts being when will this thing blow up? And these two IPOs in 2026 might have that answer. So let's look closer at this. Not the hyperbole or the marketing, but the foundations underneath. So I need to take 60 seconds to share something that might be a little boring, but is very important.

Because even though I think that 90% of the folks who watch this channel understand this, we need to demystify LLMs, large language models, and the world of artificial intelligence generally. Because without this understanding, everything else feels more magical and unknowable than it is. So feel free to skip this if you already feel like you get it. Artificial intelligence is a broad category of technologies and there is no universal definition of what qualifies. There's many kinds of artificial intelligence.

It's kind of like talking about energy. You could say coal or nuclear or solar or hydro. And then there's stuff like, you know, cold fusion that isn't real today, but still gets all sorts of hype. They are all types of energy production, but they are very, very different applications and technologies and levels of understanding for each of these. And large language models are one kind of quote artificial intelligence. And it's what's making news right now. And the way that these systems work is what's called next token prediction. Very simply, it's a giant math equation. They break words into small fragments which are called tokens. And based on huge training data, they predict how these tokens will line up in a response based on the tokens that you input to the system, meaning the questions that you ask. So, I'm going to oversimplify here, but these models learn in essentially two ways. Data and humans. Data meaning we feed these models massive and massive amounts of texts. And from all that text, they learn patterns. The second and what's becoming more important in more recent times is humans or what's called reinforcement learning from human feedback or RLHF.

So, basically, people look at the answers and then they rate them and then over and over and over again, millions of times. that teaches the model what kind of responses to give. Now there's an increasingly important movement toward self-recursive learning which is a method where these models self-improve. We will come back to that shortly because first we need to start with the strongest case for the limitless from the one person who has the most invested in it being true. So Dario Amodei runs Anthropic. He's also maybe more than anyone alive the person who has bet his entire career on this concept of scale. The way he describes scaling LLMs towards intelligence is like a chemical reaction. You take three ingredients which are data, compute and then the size of the model. And if you combine them in the right proportions, you get one thing: intelligence. And the more ingredients, you get more of the output. And for about a decade or so now, he's proved skeptics wrong again and again. So in his words, at every stage of scaling, there are always arguments and every time we manage to either find a way around or scaling is just the way around. Now the implications of this are profound. It says that if we simply make bigger and bigger models, that is essentially how we get to super intelligence. And it's important to recognize, by the way, how commercially attractive this idea is because the single thing that every investor wants is predictable growth. So if more compute equals smarter, then what he has essentially produced is an infinite money glitch. And if that's right, then maybe these valuations could be justified, right? In fact, maybe they're cheap. But is this true? Because the reality of this statement impacts all parts of your life. Every time you decide whether you're going to learn a skill or switch careers or tell your kids where to focus or you make an investment, you're making a bet. So now the loudest skeptic, his name is Ilia Sutskever and he's the co-founder of OpenAI and he's the man who arguably proved that scale worked at all in the first place. And in November of 2025, he said this. The 2010s were the age of scaling. But now that era is over and we are back in the age of research, meaning the easy gains from just building bigger and bigger are gone and what's left is hard. It's the uncertain parts that nobody has solved yet. And his reason is concrete. These models ace the tests, but quote generalize dramatically worse than people. They will pass a brutal coding exam, but then they can't fix a simple bug. It's a jagged frontier of knowledge. And it's not something that a bigger model fixes. And it's not just suits. The CEO of Google DeepMind says we need quote one or two more fundamental breakthroughs on the order of the transformer itself or AlphaGo to actually reach AGI which is a bit like saying we just need to invent time travel and then teleportation and after that it should all be pretty straightforward. So this is something that many of us intrinsically feel working with these tools. They are crazy if you stop and you look at what they do. But they also have a lot of problems and they are now getting incrementally better, not exponentially so. So if this is true and big models are not reliably getting better and quote agentic functionality can't just fix this, then the question remaining to value these companies is that of moat. Because if you're going to bet a trillion dollars on these companies that bigger is better and scale is the one thing keeping leaders ahead, then the simple math becomes how far behind is everybody else? And according to a study from MIT, the answer is that competitors are nipping at their heels. What the study shows is that open models which are free and anyone can download and run are about 90% of the performance of closed frontier models at release. and then they close most of the rest of that gap within months. And all this is at a fraction of the cost. The researchers found that the closed models cost around six times more for what they called quote only modest performance advantages. On humanity's last exam, which is one of the hardest benchmarks ever built, there were 2500 graduate level questions cited in Stanford's AI index. And then a free open Chinese model called Kimmy scores in the same tier as Claude and GPT and beats the flagship models from Meta and Amazon outright. The MIT study is from 2025, but even on brand new tests that are built specifically from problems that the models couldn't have trained on, like one called Aerospench. Early reports put open models right near

the top. And Sutzgiver even hints that in fact the generalizing of these models, meaning bigger, might actually make them worse rather than better. So, as we've discussed in previous videos, these models don't reason about your specific situation. They fall back on the most common trading patterns that they see.

When Harvard Business School asked the leading models for strategy advice across dozens of different companies, all of them got back the same fashionable answers regardless of context. And feeding in more detail, right, more context to improve the answers barely changed the result at all. And they called this trend slop. Microsoft is OpenAI's single largest backer. And this month in June of 2026, it began moving its main enterprise AI product to usage based pricing. And it started looking at free self-hosted Chinese model DeepS to run it instead of the frontier models because the cost is just too much at scale. This is literally one of the largest players in the game who has is heavily invested in open AI saying I don't think we need all this compute to make this work. This raises, by the way, a fascinating potential idea that I hope to cover in another video, which is what does a world look like where every person in the world can train their own unique models, right? Free of hyperscalers.

What if you can choose some very narrow professional specialization with data that only you have and case studies only you have and use these models to make you more and more and more valuable? How much model does this unlock for you? But for now, let's distill. There's two possible bets. The first is that massive scale makes massively more useful large language models that can't be replicated. This is not about AI as a whole. That is the question at hand. And the second is is there a fundamental flaw in the limitless analogy. So back to 2000 when Chanos exposed Enron, he did the simple thing. He looked and our question becomes is the price of these companies justified by anything real that exists today or is it based on a prophecy? And are these IPOs what leads to the pop?

One of the biggest misconceptions about technology bubbles is that they fail because the technology itself doesn't work. But in fact, bubbles burst when wealth has to be converted into money. Which brings us back to our IPOs. The moment these models stop seeming limitless and instead seem like normal, useful, and improving technology, a trillion dollar valuation is going to start to seem very far away from reality. And then the premium evaporates and everyone holding shares at once realizes the same thing. They need to sell before the other guy does. Because if these guys can't convert their wealth into money before that moment, they are left holding the bag at precisely the time when they need more money than ever to keep scaling these models the way they're being sold. But every single time a technology has been sold as limitless throughout the whole of human history. The limits have always been found. In 1982, the wires warned that personal computers would make secretaries obsolete by the end of the decade. We still have secretaries. None of that happened. The ATM was supposed to wipe out bank tellers. Instead, the number of tellers grew. The spreadsheet was supposed to end accounting. Instead, the number of accountants have roughly quadrupled. Most recently, radiology technicians were supposed to be automated away. But what we see actually is the demand for these jobs going up.

Every single time it's the same phrase. This time it's different. And every single time it hasn't been because the rhetoric of Limitless is actually limited by something nameable. Can we keep scaling large language models in a way that makes them better? Not what are the limits of AI as a category? Not are these useful tools, not do they

do helpful things? Because if it's about the utility, then there's a dollars and cents reality to the valuation. And this brings us back to David Atwater. The current hype cycle around massive job loss and society level transformation and limitless potential is all banked on a fundamental assumption. Bigger large language models are better. And right now these companies are racing towards IPO based on this story. And it's one that no data is currently showing to be true. So if the bubble bursts and they IPO first, it's not those guys who are going to be left holding the bag. It's us. And what I'm imploring for all of us is take the time to figure out what you think is true. Because if the prophecy is these large language models can do our thinking for us, then the most valuable skill is the ability to think clearly and rationally for ourselves. So where does this leave us? I can't tell you when this pops. No one can. When channels looked at Enron, he never predicted a date. What he did was understand the value at its foundations and bet accordingly. So let's do the same thing. There are three things that decide this. The first is does the scale story hold? If the next generation of models is a genuine leap, then maybe this is right. But if they are incrementally better and not magic, it's prophecy. Two, does the moat hold? So watch the free and open models. if they stay within months of the frontier at a fraction of the price. There is no moat.

And three, can the actual businesses support the checks that they're writing? Chanos is pointing at the debt these companies took on against chips that lose value faster than the loans assume. And if those defaults start, then prick, that's the pin. So now believers in this whole idea will have one last rebuttal. And they'll say that even if the models stop improving, that doesn't mean scaling is dead. What it means is we're measuring the wrong thing. Once we fix that, the gains come back. Meaning that they can keep learning and improving themselves. But notice what that argument does.

It's impossible to ever lose because it's always working. And if it looks like it's not, it's just it hasn't come yet. Those who live by the crystal ball are destined to eat broken glass. And that works if you're a venture investor who is in early and bought this on pennies for of pennies of the dollar. It doesn't work if you're the one buying late cycle holding the bag. So I implore you to do one simple thing. Look and then think for yourself what here is true. And now you've seen through it.

So, if you want the story of how the founders are racing to cash out before any of this lands, watch the IPO video next. And if you want to understand why I feel like these models are actually in many cases delivering negative, not positive value, see the billionaires ne video next. And if you've seen all those, I'll leave a few others in the description.