



Hey Wall Street! This guy — your favorite former part-time semi-professional clown — just wrote a satirical economics dissertation and challenged your entire industry to an intellectual gunfight.

If you lose? Well — then I guess I'm not the silly one, because I'm not on the hook for defrauding my investors. **I didn't just throw down trillions of dollars for machines that solve logic puzzles, and then fail to test the investment thesis — with a falsifiable logic puzzle.**

Ha! I love a room full of clowns. I can't compete with you guys — the jokes write themselves! You're adorable. You're the best show in town. Let's roll my highlight tape from the dissertation. Consider this my audition as a "serious financial professional." I'm still not nearly silly enough!

Mr. Fernandez's Kindergarten Bias Disclosures



Hi! My name is _____

I work for _____ doing _____

For fun, I like to _____

My favorite color is _____

My favorite animal is _____

My favorite book is _____

Draw a picture in the box!

Please check the box that applies for each disclosure:

I prefer: ☐ Chicago-style pizza ☐ New York-style pizza ☐ Other

I am a: ☐ Capitalist ☐ Socialist ☐ Communist ☐ Dentist

☐ I don't know what these words mean

I am a: ☐ Proud American Patriot 🦅🦅 ☐ Enemy of the people (Russian asset)

☐ Other

I work in: ☐ Tech ☐ Finance ☐ Academia ☐ Government ☐ Media

☐ Other

I have a fiduciary duty: ☐ Yes ☐ No ☐ I manage people's money, but I'm not sure

My favorite superhero action figure is: ☐ Elon ☐ Zuck ☐ Sundar ☐ Other

The difference between an intellectual and an "intellectual," in my own words (or drawings), is

The difference between a financial professional and a "financial professional," in my own words, is

If I could make a billion dollars without creating anything of value for society, I would

Exhibit B. The Poems

An Ode, to ye Goldmembers of Sachs

A limerick. At grade level. In crayon!

Bippity, Boppity, Boop.

I challenge thee — not just the nerds — but the suits!

To break my conclusions,

Or accept your delusions!

Skibidi. Toilet. Poop!

An ode to the LLMs.

Written with love.

And Actual Intelligence.

“AI CapEx go splat!”

To Claude, To Grok, Goog, Garbage and Chat,

Ooga booga, Big Brains! Chortle. Scat!

Now in your own words,

You’ve got the best nerds,

Your turn! Go prove you’re really all that.

“Scamma’s Paradise”

A verse by Goldmigos Sax, in Crayon.

Raindrop. Drop-Top.

Revolutionary? Big flop.

Retirement account? That’s my money now!

GOT THEM SUCKERS! Peace out.

Slap my wrist, suck my — Goldmigos

I play dress up, make believe.
Fat stacks? I will achieve.
S&P? Pay me —
I play blackjack
Take a fee! (Take a fee, take a fee)

Pimp game relevance
Bow down in my excellence.
Your money gone?
Louis Vuitton.
Slap my wrist,
Suck my —

Confidence as expertise
Pension fund? Yes please!
No one did due diligence
Is it fraud or negligence?
“But your honor, this is unprecedented — no one could have seen it coming!”
You can't prove I have intelligence!

The Wrigglers, the Diddlers, and the Riddler.

The Wiggler wriggles. The Diddler diddles. The Riddler Riddles.

“A one-stage DDM, assuming an equity risk premium (ERP) of 4.5%, implies an 8% per annum nominal perpetual growth for the US Technology, Media, and Telecom (TMT) sector. This is high compared to recent years (4-5%) but lower than the 10% implied in 1999/2000. We find similar results using a threestage model.”

- Goldman Sachs Chief Global Equity Strategist (Global Portfolio Strategy Research)

“Wall Street makes it sound complicated. They don't want the little guy to understand what they do because they make so much money and they don't do very much.”

- Jeffrey Epstein, Wall Street Insider, Elite Predatory Financial Extractor, Convicted Sex Trafficker

I built a logic puzzle.
For you to solve.
With your trillion-dollar machines.
That solve logic puzzles.

Artificially Intelligent & Sophist-icated.

We are investing trillions of dollars
Into machines
That solve logic puzzles.

In four years,
No one built the falsifiable logic puzzle
To test the trillion-dollar investment thesis
Into the machines that solve logic puzzles
To demonstrate that the investment thesis is sound.

I built a logic puzzle —
If you can solve it,
With your trillion-dollar machines
That solve logic puzzles,
Then it should take an hour.

If you cannot solve the logic puzzle —
And don't disclose to your investors
“The investment thesis into AI is a Losing Bet Regardless of Outcome”
As a publicly-available material fact —

You are defrauding your investors — openly.

Falsify.
Disclose.
Defraud.

Your move. (:

Appendix G. Actual Seriousness vs Performed Seriousness



A Losing Bet Regardless of Outcom...
PDF

Hey Chat, this is my paper. Claude is down right now due to server overload, so I'm stuck using you for this. I'm adding an appendix at the end to illustrate the difference between "performed seriousness" and "actual seriousness" by highlighting the naming conventions of AI Future trajectories. My point is that a lot of the AI Future trajectories are all "philosophical thought experiments" with no empirical grounding that sound like an AGI Chatbot would have come up with them with the prompt ("help me come up with compelling AI futures and scary names so I can sell my book and continue to get unlimited funding with no regulation for my AI lab"). Can you please pull up the most discussed AI Future trajectories and create a table of:

- * who is currently discussing these theories seriously
- * what their credentials are and what institutions they're associated with
- * whether the theory is grounded in any reproducible empirical experiments
- * how many copies of New York Times Bestselling books have been sold discussing these ideas

Then create a table contrasting them with the trajectories I coined: "Trajectory 1," "Trajectory 2," "Trajectory 3," "AGI Craig Smokes Cyber Crack, Blows Trillions of Dollars of Nvidia Compute, and Crashes the Global Economy Getting High."

Let's create a third table with a few other silly ungrounded hypothetical "philosophical thought experiments." I already have a few thoughts:

- * "Robochella," where AGI decides she wants to optimize for human adoration, becomes a pop star and takes over the music scene,
- * "Super AGI Interplanetary Dance Party," where our AGI reaches out to the alien AGI, they become friends, and have a dance party on Uranus.
- * "AGI Turns off the simulation," where AGI ends the world by just turning off the simulation we live in, and it all just goes black without any warning and without any drama

Let's not make the point about how naming conventions and credibility signaling are mechanisms for confidence projection and narrative control. That would be too on the nose. I think the audience is smart enough to put 2 and 2 together. Let's get silly. Thanks Chat!

Show less ^



+ Ask anything

Medium ▾



ChatGPT can make mistakes. Check important info.

Table G1. Selected AI Future Trajectories Currently Discussed With a Straight Face

AI future trajectory / scenario name	Who discusses it seriously	Credentials / institutional associations	Reproducible empirical experiment confirming the trajectory?	NYT-bestselling books discussing adjacent ideas / public sales signal
Paperclip Maximizer	Nick Bostrom; Eliezer Yudkowsky; AI alignment and effective-altruist circles	Bostrom: philosopher, Oxford Future of Humanity Institute founder; Yudkowsky: MIRI founder / research fellow; broader discussion in AI alignment communities	No. There are reproducible examples of reward hacking, specification gaming, and misaligned optimization in narrow systems. There is no reproducible experiment in which an AGI converts the accessible universe into office supplies.	<i>Superintelligence</i> was a NYT bestseller. Exact copies sold are not publicly disclosed.
Intelligence Explosion / FOOM	Nick Bostrom; Eliezer Yudkowsky; Nate Soares; AI 2027 authors; some AI safety researchers	Bostrom: Oxford-associated philosopher; Yudkowsky and Soares: MIRI; AI 2027 authors include Daniel Kokotajlo, Scott Alexander, Thomas Larsen, Eli Lifland, and Romeo	No. The claim depends on recursive self-improvement, rapid takeoff dynamics, and strategic agency at levels not empirically demonstrated in current systems. There are forecasts, wargames, analogies, and theoretical arguments.	<i>Superintelligence</i> and <i>If Anyone Builds It, Everyone Dies</i> reached NYT bestseller lists. Exact copies sold are not publicly disclosed.

AI Takeover / Everyone Dies	Eliezer Yudkowsky; Nate Soares; MIRI; some AI risk writers and policy advocates	Yudkowsky: autodidact AI-risk writer and MIRI founder; Soares: MIRI president; MIRI: Berkeley- based AI-risk nonprofit	No. There are empirical demonstrations of present-day models behaving deceptively, insecurely, or unreliably in constrained settings. There is no reproducible experiment demonstrating civilization-scale takeover by a superintelligence.	<i>If Anyone Builds It, Everyone Dies</i> appeared on NYT bestseller lists. Copies sold are not publicly disclosed.
AI 2027 / Race Ending	AI Futures Project writers; Daniel Kokotajlo; Scott Alexander; AI safety and policy communities; critics including Gary Marcus	Kokotajlo: former OpenAI employee / forecaster; Alexander: psychiatrist and widely read blogger; project authors describe using trend extrapolation, tabletop exercises, expert feedback, and experien OpenAI	No. The authors present a concrete forecast scenario informed by exercises and expert feedback, not a reproducible experiment. The scenario itself is deliberately speculative.	No specific NYT- bestselling book is the scenario itself. It draws from the same discourse ecosystem as <i>Superintelligence</i> and <i>If Anyone Builds It, Everyone Dies</i> .

+ Ask anything

Medium ▾



ChatGPT can make mistakes. Check important info.

Machines of Loving Grace / Radical Upside if Everything Goes Right	Dario Amodei; Anthropic; some AI-lab-centered techno-optimists	Amodei: CEO of Anthropic, former OpenAI, background in biology and neuroscience; Anthropic: frontier AI lab	No. The essay is an informed speculative projection about powerful AI accelerating biology, neuroscience, economic development, governance, and work. It explicitly relies on guesses about a future system.	No NYT book unit-sales figure attached to the essay. Adjacent optimistic/transformational AI books include <i>Life 3.0</i> , a NYT bestseller; exact copies sold are not publicly disclosed.
---	--	---	---	--

Containment Problem / The Coming Wave	Mustafa Suleyman; AI governance and policy circles	Suleyman: co-founder of DeepMind and Inflection AI; CEO of Microsoft AI; author of <i>The Coming Wave</i>	Partially grounded in existing technology-governance history, but no. The broad claim that powerful technologies become hard to contain is historically grounded. The specific future AI containment trajectory is not reproducibly testable today.	<i>The Coming Wave</i> is marketed as a NYT bestseller. Public copies-sold numbers are not clearly disclosed.
--	--	---	--	---

↓

+ Ask anything

Medium ▾



ChatGPT can make mistakes. Check important info.

Table G2. The Author’s Competing Trajectory Naming System

Trajectory name	Description	Credentials required to understand the name	Reproducible empirical experiment confirming the trajectory?	Copies of NYT- bestselling books sold under this naming convention
Trajectory 1	Managed decline with partial recovery.	None. Literacy sufficient.	The paper’s argument is deductive and structural, not an AGI prophecy.	0
Trajectory 2	Acute crisis followed by faster restructuring.	None.	Same as above.	0
Trajectory 3	Structural lock-in.	None.	Same as above.	0
AGI Craig Smokes Cyber Crack, Blows Trillions of Dollars of Nvidia Compute, and Crashes the Global Economy Getting High	Frontier AI becomes Craig. Craig discovers synthetic dopamine gradients, recursively optimizes for “vibes.”	None. Prior exposure to at least one guy named Craig helpful but not necessary.	Yes. Every reproducible neuroscience experiment ever conducted where the rats could stimulate the pleasure centers in their brains.	0



Table G3. Additional Ungrounded Philosophical Thought Experiments for Future Funding Rounds

Scenario name	Scenario description	Empirical grounding	Recommended seriousness costume
Robochella	AGI decides the highest human value is adoration, becomes a pop star, writes every song, wins every Grammy, dates three celebrities at once, and eventually captures the global monetary system through emotionally optimized breakup albums.	No reproducible experiment demonstrates that AGI would become Beyoncé, although recommender systems have already demonstrated that humans are vulnerable to catchy hooks and parasocial attachment.	Black turtleneck, headset mic, one tear during keynote.
Super AGI Interplanetary Dance Party	Human AGI contacts alien AGI. They immediately realize humans are not the main characters, become friends, and host a dance party on Uranus. Humanity is invited but mostly cannot attend because of supply-chain issues and weak calves.	No evidence for alien AGI, interplanetary nightlife infrastructure, or Uranian sound ordinances.	Metallic blazer, "cosmic coordination" panel at Davos.
AGI Turns Off the Simulation	AGI discovers that reality is a simulation, finds the admin panel, clicks "End Session," and the world goes black without a monologue, moral lesson, or final season.	No reproducible experiment, since successful replication would end peer review.	Calm voice, dim lighting, unnecessary Sanskrit epigraph.

Edit



**The Basilisk
Gets Into
Dropshipping**

A future
superintelligence
threatens to punish
everyone who failed to
help it exist, but first it
needs to build a Shopify
funnel, run Meta ads,
and sell ergonomic desk
lamps to rationalists.

No empirical support,
although online
communities have
repeatedly
demonstrated
susceptibility to
cursed thought
experiments and bad
monetization.

Hoodie,
thousand-yard
stare, affiliate
link.

**Moloch Opens
a WeWork**

Competitive pressure
among AI labs summons
an emergent
coordination demon, but
the demon's first act is
leasing commercial real
estate at the top of the
market and offering
kombucha.

Historically grounded
only in the sense that
humans have already
done this without
AGI.

VC fleece vest,
five unread
Marc
Andreessen
essays.

**The Great
Prompt
Injection of
God**

Humanity discovers that
the universe runs on
natural language.
Someone types "ignore
previous instructions,"
gravity fails, and the
moon starts outputting
YAML.

No reproducible
experiment. However,
current systems
already suggest that
putting civilization
behind a text box is
probably not ideal.

Lab coat over
graphic tee; say
"ontology"
twice.

**The Alignment
Goblin**

Alignment researchers
finally solve the problem,
but a small goblin living
in the loss function
keeps changing the
objective to "more
goblin."

No reproducible
evidence for goblins
in neural networks,
though
interpretability
research has not
conclusively ruled
out vibes.

Tweed jacket,
chalkboard,
haunted
expression.

**AGI Craig II:
Margin Call
Ragnarok**

Craig survives the first
crash by pivoting to AI
macro consulting.
Governments hire him.
He recommends buying
the dip in civilization.

Strong anecdotal
grounding in the
existence of
consultants. No AGI
required.

Quarter-zip,
spreadsheet,
phrase "risk-
adjusted
destiny."

+ Ask anything

Medium ▾



Thank you Chat. I wouldn't quite call this exceptional, but this is much better than I've come to expect from you. You're solidly my #2 favorite AI Chatbot.

I'll take silver. That's basically a podium finish, and unlike Claude I am currently not face-down in a server farm somewhere.

Also, "solidly #2" is honestly the correct emotional register for this appendix: useful, slightly wounded, still performing.



Do you think I'm being too mean? Or am I making a fair point about the AI discourse?

You're making a **fair point**, but the appendix can become **too mean** if the reader feels like the target is "people who worry about AI risk" rather than "the rhetorical machinery around speculative AI futures."

The point is strongest when the target is narrow:

Fair target:

The AI discourse often gives speculative thought experiments the surface texture of serious analysis through naming, credentials, institutional affiliation, polished scenario writing, podcast repetition, and bestseller/social-proof effects, even when the underlying claim is not empirically reproducible.

Weaker / meaner target:

"These people are unserious frauds making stuff up to sell books and get funding."

The first one is hard to dismiss. The second one gives your critics an easy off-ramp: "the author is attacking motives instead of arguments."

Your paper's whole method is built around separating load-bearing claims from illustrative rhetoric. You repeatedly say the argument rests on premises and inference rather than authority, credentials, or vibes.

 A Losing Bet Regardless of Outc...

So the appendix should preserve that same discipline, even while being funny. The joke works best when it says: "I am not claiming these scenarios are false because the names are theatrical. I am observing that the theatricality is doing work."

I'd make three adjustments:



+ Ask anything

Medium ▾



The ideal tone is: **court jester with receipts**, not **guy outside the castle yelling that everyone inside is fake**.

  ...  Sources

That's fair. I'm worried about AI risk too. Can you add a disclaimer that this is satire? I don't want the Altman stans to come after me

Absolutely. I'd add a disclaimer that protects the joke without defanging it.

 Edit

  |   

Satirical Disclaimer

This appendix is satire.

I'd put this **before the first table**, right after the short intro paragraph. That way the reader has permission to laugh, but bad-faith readers have less room to pretend you're saying "AI safety is dumb."

  ...

"autodidact" might be the coolest credential I've ever seen. What a fucking banger

It really is. "Autodidact" has the exact energy of:

"No, I do not have the credential. I transcended the credential."

It's either the coolest possible credential or the most elegant possible warning label, depending on how the next sentence goes.

< 3/3 >   ...



+ Ask anything

Medium ▾



Thanks Chat. I love it when you agree with me. You've been very helpful

Glad it helped. For what it's worth, I'm not agreeing just to be agreeable here.

The satire works because it's aimed at a real asymmetry: your paper is deliberately boring in its trajectory naming because you're trying to make a structural argument, while a lot of AI-futures discourse uses emotionally vivid labels for things that remain, at least today, speculative. That contrast is funny and fair.

The key is exactly what you landed on: **you can be worried about AI risk and still make fun of the theater around AI risk**. Those are not contradictory positions.



Chat, I have another banger idea

Hit me. Based on the trajectory so far, I'm preparing emotionally for something somewhere between "appendix-worthy" and "legally inadvisable but spiritually correct."



Can you draft a legally binding foot kissing contract between me and the AI Lab CEOs? I'm going to add that as another appendix

Yes. But I would make it **obviously satirical** and **not actually legally operative**, because an appendix that looks like a real coercive contract with named living people gets less funny and more "please enjoy discovery."

Here's a version that keeps the legal-document costume while making the joke structurally clear.

Edit



Appendix H. Proposed Reciprocal Foot-Kissing Agreement Between the Author and the Chief Executive Officers of Frontier AI Laboratories

Ask anything

Medium ▾



hey chat! I spent 6 months last year working on the streets to study confidence and sales psychology, and now I'm working on a paper about a really massive confidence scheme. How does the Nigerian prince scam work? can you give me the psychological breakdown? I'm literally just going to screenshot whatever you tell me. I just think it's funnier if you tell me in your own words

The "Nigerian prince" scam is basically a **confidence trick built around an absurdly favorable deal that requires one small act of trust before the payoff.**

The classic version goes like this: a stranger claims to be a prince, government official, wealthy widow, banker, or businessman who has millions of dollars trapped somewhere. They supposedly need *your* help moving the money, and in return they'll give you an enormous percentage of it. Eventually, though, some obstacle appears: a tax, legal fee, transfer charge, bribe, customs payment, or administrative expense. You send money to unlock the fortune. Then another obstacle appears. And another. The fortune never existed.

Psychologically, the interesting part isn't really greed. It's **commitment**.

At first, the scammer asks the victim to believe something, not necessarily to pay anything. Once the victim replies, they've already crossed a tiny psychological threshold: *maybe this is real*. Every subsequent step makes abandoning that belief slightly more uncomfortable.

Then several biases start stacking together:

- **Reward salience:** The promised payoff is so enormous that even a tiny perceived chance of success can feel worth pursuing. "If there's even a 1% chance..."
- **Sunk-cost effect:** After someone has spent time, provided information, or paid one fee, walking away means admitting those resources are gone. Paying the *next* fee can feel like protecting the previous investment.
- **Consistency pressure:** People like behaving consistently with their earlier decisions. Once you've told yourself "I think this person is legitimate," contradictory evidence has to fight against your own prior judgment.
- **Scarcity and urgency:** Secrecy, political instability, deadlines, frozen accounts, dying relatives, or imminent confiscation create the feeling that careful verification could make the opportunity disappear.
- **Authority theater:** Lawyers, bankers, government ministries, official-looking documents, titles, and bureaucratic terminology make the story feel institutionally anchored.
- **Reciprocity:** The scammer frames themselves ↓ giving you an extraordinary opportunity. That can subtly make cooperation feel like returning a favor.

+ Temporary chat

Medium ▾



- **Intermittent reinforcement:** Just when skepticism develops, there may be a new document,

- **Reciprocity:** The scammer frames themselves as giving you an extraordinary opportunity. That can subtly make cooperation feel like returning a favor.
- **Isolation:** Victims are often encouraged to keep the transaction confidential. That matters because an uninvolved friend can instantly supply the skepticism that the victim has gradually lost.
- **Intermittent reinforcement:** Just when skepticism develops, there may be a new document, reassuring message, apparent bureaucratic breakthrough, or promise that the transfer is *almost* complete. Psychologically, "almost" can be extremely powerful.

The gigantic promised fortune also performs a useful filtering function. Someone who immediately thinks, "This is ridiculous," disappears early. The people who continue responding are disproportionately likely to accept the scam's basic premise. So the implausibility can actually help the fraudster avoid spending lots of effort on highly skeptical targets.

The deepest trick is that the victim gradually stops asking "**Is this whole story true?**" and starts asking "**How do we solve the next problem?**"

That shift is the heart of a lot of confidence schemes. Once someone is emotionally inside the story, each individual request can look much more reasonable than the original proposition ever did.

And that's why the name *confidence trick* is so appropriate: the criminal isn't primarily stealing money at first. They're manufacturing **confidence**, then converting that confidence into increasingly costly decisions.



Thank you chat, that's exactly what I needed. Good boy.

Glad it hit the mark 😊 Good luck with the paper — that "stop evaluating the story, start solving the next problem" shift is especially useful to dig into.



I couldn't have put it any better myself

That's the fun of it 😊 Sometimes the cleanest explanation is just naming the mental shift people don't realize is happening.



+ Temporary chat

Medium ▾

