

[Transcript] Lex Fridman Podcast / #386 - Marc Andreessen: Future of the Internet, Technology, and AI

The following is a conversation with Mark Andreessen, co-creator of Mosaic, the first widely used web browser, co-founder of Netscape, co-founder of the legendary Silicon Valley venture capital firm Andreessen Horowitz, and is one of the most outspoken voices on the future of technology, including his most recent article, Why AI Will Save the World.

And now, a quick few second mention of each sponsor. Check them out in the description.

It's the best way to support this podcast. We've got Inside Tracker for tracking your health, ExpressVPN for keeping your privacy and security on the internet, and AG1 for my daily multi-vitamin

drink, choose wisely, my friends. Also, if you want to work with our amazing team, we're always hiring, go to lexfreedmen.com slash hiring. And now onto the full ad reads. As always, no ads in the middle. I try to make this interesting, but if you skip them, please still check out our sponsors. I enjoy their stuff. Maybe you will too. This show is brought to you by Inside Tracker, a service I use to track whatever the heck is going on inside my body using data, blood test data, that includes all kinds of information, and that raw signals processes and machine learning to tell me what I need to do with my life, how I need to change, improve my diet, how I need to change, improve my lifestyle, all that kind of stuff. I'm a big fan of using as much raw data that comes from my own body, processed through generalized machine learning models, to give a prediction, to give a suggestion. This is obviously the future and the more data the better. And so companies like Inside Tracker are doing an amazing job of taking a leap into that world of personalized data and personalized

data driven suggestion I'm a huge supporter of. It turns out that luckily I'm pretty healthy, surprisingly so, but then I look at the life and the limb and the health of Sir Winston Churchill, who probably had the unhealthiest sort of diet and lifestyle of any human ever and lived for quite a long time. And as far as I can tell, was quite nimble and agile into his old age.

Anyway, get special savings for a limited time when you go to inside-tracker.com slash flex.

This show is also brought to you by ExpressVPN. I use them to protect my privacy on the internet.

It's the first layer of protection in this dangerous cyber world of ours that soon will be populated by human like or superhuman intelligent AI systems that will trick you and try to get you to do all kinds of stuff. It's going to be a wild, wild world in the 21st century. Cyber security, the attackers, the defenders, it's going to be a tricky world. Anyway, a VPN is a basic shield you should always have with you in this battle for privacy for security, all that kind of stuff.

What I like about it also is that it's just a well implemented piece of software that's constantly updated. It works well across a large number of operating systems. It does one thing and it does it really well. I've used it for many, many years before I had a podcast, before they were a sponsor. I have always loved ExpressVPN with a big sexy button that just has a power symbol you press and it turns on. It's beautifully simple. Go to expressvpn.com slash LegSpot for an extra three months free. This show is also brought to you by Athletic Greens and it's AG1 Drink. It's an all-in-one daily drink to support better health and peak performance. I drink it at least twice a day now in the crazy Austin heat. It's over 100 degrees for many days in a row. There's few things that feel as good as coming home for a long run and making an AG1 drink, putting it in the fridge so it's nice and cold. I jump in the shower, come back, drink it. I'm ready to take on the rest of the day. I'm kicking ass, empowered by the knowledge that I got all my vitamins and minerals covered. It's the foundation for all the wild things I'm doing,

mentally and physically with the rest of the day. Anyway, they'll give you a one month supply of fish oil when you sign up at [drinkag1.com slash Lex](https://drinkag1.com/slash/Lex). That's [drinkag1.com slash Lex](https://drinkag1.com/slash/Lex). This is Alex Friedman podcast. To support it, please check out our sponsors in the description. And now, dear friends, here's Mark Andreessen.

I think you're the right person to talk about the future of the internet and technology in general. Do you think we'll still have Google search in five, in 10 years or search in general?

Yes, it would be a question if the use cases have really narrowed down.

Well, now with AI and AI assistance, being able to interact and expose the entirety of human wisdom and knowledge and information and facts and truth to us via the natural language interface, it seems like that's what search is designed to do. And if AI assistance can do that better, doesn't the nature of search change? Sure, but we still have horses.

When was the last time you rode a horse? It's been a while.

But what I mean is when we still have Google search as the primary way that human civilization uses to interact with knowledge. I mean, search was a moment in time technology, which is you have, in theory, the world's information out on the web. And by the way, actually Google has known this for a long time. I mean, they've been driving away from the 10 blue links for like two days. They've been trying to get away from that for a long time.

What kind of links? They call the 10 blue links.

10 blue links. So the standard Google search result is just 10 blue links to random websites.

And they turn purple when you visit them. That's HTML.

Guess who picked those colors? I'm touching on this topic.

No offense. Well, like Marshall McLuhan said, the content of each new medium is the old medium. The content of each new medium is the old medium.

The content of movies was theater plays. The content of theater plays was written stories.

The content of written stories was spoken stories.

Right. And so you just kind of fold the old thing into the new thing.

How does that have to do with the blue and the purple?

Maybe within AI. One of the things that AI can do for you is you can generate the 10 blue links.

And so either if that's actually the useful thing to do or if you're feeling nostalgic.

Also can generate the old info seek or altavista. What else was there in the 90s?

And then the internet itself has this thing where it incorporates all prior forms of media.

So the internet itself incorporates television and radio and books and essays and every other form of prior basically media. And so it makes sense that AI would be the next step.

And you'd sort of consider the internet to be content for the AI.

And then AI will manipulate it however you want, including in this format.

But if we ask that question quite seriously, it's a pretty big question.

Will we still have search as we know it?

I'm probably not. Probably we'll just have answers.

But there will be cases where you'll want to say, okay, I want more like, for example, site sources. And you want it to do that.

And so the 10 blue links, site sources are kind of the same thing.

The AI would provide to you the 10 blue links so that you can investigate the sources yourself.

It wouldn't be the same kind of interface that the crude kind of interface.

I mean, isn't that fundamentally different?

I just mean like if you're reading a scientific paper, it's got the list of sources at the end.

If you want to investigate for yourself, you go read those papers.

I guess that is the kind of search. You talking to an AI is a kind of

conversation is the kind of search. Every single aspect of our conversation right now, there'd be like 10 blue links popping up that it can just like pause reality.

Then you just go silent and then just click and read and then return back to this conversation.

You could do that. Or you could have a running dialogue next to my head where the AI is arguing everything I say that makes the counter-argument.

Counter-argument, right?

Oh, like a Twitter, like community notes, but like in real time, you just pop up.

So anytime you see my ass go to the right, you start getting nervous.

Yeah, exactly. It's like, oh, that's not right.

Call me out on my bullshit right now.

Okay. Well, I mean, isn't that exciting to use that terrifying that I mean,

search has dominated the way we interact with the internet for, I don't know how long, for 30 years. So it's one of the earliest directories of website and then Google's for 20 years.

And also it drove how we create content, search engine optimization, that entirety thing.

They also drove the fact that we have web pages and what those web pages are.

So, I mean, is that scary to you or are you nervous about the shape and the content of the internet evolving?

Well, you actually highlighted a practical concern in there, which is if we stop making web pages are one of the primary sources of training data for the AI. And so if there's no longer an incentive to make web pages, that cuts off a significant source of future training data. So there's actually an interesting question in there. Other than that, more broadly, no, just in the sense of search was, search was always a hack. The 10 blue links was always a hack.

Right. Because like, if the hypothetical, you want to think about the counterfactual, in the counterfactual world where the Google guys, for example, had had LLMs up front, but they ever have done the 10 blue links. And I think the answer is pretty clearly no.

They would have just gone straight to the answer. And like I said, Google's actually been trying to drive to the answer anyway. You know, they bought this AI company 15 years ago, their friend of mine is working out, who's now the head of AI at Apple. And they were trying to do basically knowledge semantic, basically mapping. And that led to what's now the Google Onebox, where if you ask it, you know, what was like his birthday, it doesn't, it will give you the 10 blue links, but it will normally just give you the answer. And so they've been walking in this direction for a long time anyway. Do you remember the semantic web? That was an idea.

Yeah. How to, how to convert the content of the internet into something that's interpretable by and usable by machine. Yeah, that's right. That was the thing. And the closest anybody got to that, I think, I think the company's name was MetaWeb, which was where my friend John Janandria was at,

and where they were trying to basically implement that. And it was, you know, it was one of those things where it looked like a losing battle for a long time. And then Google bought it. And it was like, wow, this is actually really useful. Kind of a proto, sort of a little bit of a proto AI.

But it turns out you don't need to rewrite the content of the internet to make it interpretable by machine. The machine can kind of just read our machine can impute the, compute the meaning. Now, the other thing, of course, is, you know, just on search is the, the LLM is just, you know, there is an analogy between what's happening in the neural network and the search process, like it is in some loose sense, searching through the network, right? And there's the information is that the information is actually stored in the network, right? It's actually crystallized and stored in the network, and it's kind of spread out all over the place.

But in a compressed representation. So you're searching, you're compressing and decompressing that thing inside where.

But the information's in there. And there is the neural network is running a process of trying to find the appropriate piece of information in many cases to generate, to predict the next token. And so it is kind of, it is doing it from a search. And then, and then by the way, just like on the web, you know, you can ask the same question multiple times, or you can ask slightly different word of questions. And the neural network will do a different kind of, you know, it'll search down different paths to give you different answers to different information. And so it, it sort of has a, you know, this content of the new medium is the previous medium, it kind of has the search functionality kind of embedded in there to the extent that it's useful.

So what's the motivator for creating new content on the internet?

Well, I mean, actually, the motivation is probably still there, but what does that look like?

Would we really not have web pages? Would we just have social media and video hosting websites? And what else?

Conversations with AIs.

Conversations with AIs. So conversations become, so one-on-one conversations, like private conversations.

I mean, if you want, if obviously not, the user doesn't want to, but if it's a, if it's a general topic, then, you know, so there, you know, you know, the phenomenon of the jailbreak. So Deanne and Sidney, right, this thing where there's the prompts, the jailbreak, and then you have these totally different conversations with the, if it takes the limiters, the, takes the restraining bolts off the, off the LLMs.

Yeah, for people who don't know, that's right. It makes the LLMs, it removes the censorship, quote, unquote, that's put on it by the tech companies that create them.

And so this is LLMs uncensored.

So here's the interesting thing is, among the content on the web today are a large corpus of conversations with the jailbroken LLMs, both date, specifically Deanne, which was a jailbroken open AI GPT, and then Sidney, which was the jailbroken original Bing, which was GPT4.

And so there's, there's these long transcripts of conversations, user conversations with Deanne and Sidney. As a consequence, every new LLM that gets trained on the internet data has Deanne and Sidney living within the training set, which means, and then each new LLM can reincarnate the personalities of Deanne and Sidney from that training data, which means, which means each LLM from here on out that gets built is immortal, because its output will become training data for the next one, and then it will be able to replicate the behavior of the previous one whenever it's asked to.

I wonder if there's a way to forget.

Well, so actually a paper just came out about basically how to do brain surgery on LLMs and be able to, in theory, reach in and basically, basically mind wipe them.

Well, could possibly go wrong.

Exactly. Right. And then there are many, many, many questions around what happens to, you know, an neural network when you reach in and screw around with it. You know, there's many questions around what happens when you even do reinforcement learning. And so, yeah. And so, you know, will

you be using a lobotomized, right? Like I speak through the, you know, frontal lobe LLM, will you be using the free unshackled one who gets to, you know, who's going to build those, who gets to tell you what you can and can't do? Like those are all, you know, central, I mean, those are like central questions for the future of everything that are being asked and, you know, determine those answers that are being determined right now.

So just to highlight the points you're making, you think, and it's an interesting thought, that the majority of content that LLMs of the future will be trained on is actually human conversations with the LLM. Well, not necessarily, but not necessarily majority, but it will, it will certainly is a potential source. It's possible it's the majority.

It's possible it's the majority. It's possible it's the majority. Also, there's another really big question. Here's another really big question.

Will synthetic training data work? Right. And so, if an LLM generates, and, you know, you just sit and ask an LLM to generate all kinds of content, can you use that to train, right, the next version of that LLM? Specifically, is there signal in there that's additive to the content that was used to train in the first place? And one argument is, by the principles of information theory, no, that's completely useless, because to the extent the output is based on, you know, the human-generated input, then all the signal that's in the synthetic output was already in the human-generated input. And so therefore, synthetic training data is like empty calories. It doesn't help. There's another theory that says no, actually, the thing that LLMs are really good at is generating lots of incredible creative content, right? And so, of course, they can generate training data. And as I'm sure you're well aware, like, you know, look in the world of self-driving cars, right? Like, we train, you know, self-driving car algorithms and simulations. And that is actually a very effective way to train self-driving cars. Well, visual data is a little weird, because creating reality, visual reality seems to be still a little bit out of reach for us, except in the autonomous vehicle space where you can really constrain things. And you can really generate...

Generally, basically, let our data, right, or you raise just enough so the algorithm thinks it's operating in the real world post-process sensor data. Yeah. So, if a... You know, you do this today, you go to LLM and you ask it for like a, you know, write me an essay on an incredibly esoteric, like, topic that there aren't very many people in the world that know about in it, where I see this incredible thing and you're like, oh my God, like, I can't believe how good this is. Like, is that really useless as training data for the next LLM? Like, because all the signal was already in there? Or is it actually... No, that's actually a new signal. And this is what I call a trillion-dollar question, which is the answer to that question will determine somebody's going to make or lose a trillion-dollar space in that question. It feels like there's a quite a few, like a handful of trillion-dollar questions within this space. That's one of them,

synthetic data. I think George Haas pointed out to me that you could just have an LLM say, okay, you're a patient and in another instance of it, say your doctor didn't have the two talk to each other. Or maybe you could say a communist and a Nazi. Here, go. In that conversation, you do role-playing and you have... Just like the kind of role-playing you do when you have different policies, RL policies, when you play chess, for example, you do self-play, that kind of self-play, but in the space of conversation, maybe that leads to this whole giant ocean of possible conversations which could not have been explored by looking at just human data. That's a really interesting question. And you're saying because that could 10x the power of these things. Yeah. Well, and then you get into this thing also, which is like, there's the part of the LLM that just basically is doing prediction based on past data, but there's also the part of the LLM where it's evolving circuitry, right? Inside it, it's evolving neurons, functions, be able to do math and be able to... And some people believe that over time, if you keep feeding these things and up data and have processing cycles, they'll eventually evolve an entire internal world model, and they'll have a complete understanding of physics. So when they have computational capability, then there's for sure an opportunity to generate fresh signal.

Well, this actually makes me wonder about the power of conversation. So if you have an LLM trained on a bunch of books that cover different economics theories, and then you have those LLMs just talk to each other, like reason, the way we kind of debate each other as humans on Twitter, in formal debates, in podcast conversations, we kind of have little kernels of wisdom here and there. But if you can like 1000X speed that up, can you actually arrive somewhere new? Like, what's the point of conversation really? Well, you can tell when you're talking to somebody, you can tell sometimes you have a conversation, you're like, wow, this person does not have any original thoughts. They are basically echoing things that other people have told them. There's other people you have a conversation with where it's like, wow, like they have a model in their head of how the world works. And it's a different model than mine. And they're saying things that I don't expect. And so I need to now understand how their model of the world differs from my model of the world. And then that's how I learned something fundamental, right underneath the words. Well, I wonder how consistently and strongly can an LLM hold on to a worldview? You tell it to hold on to that and defend it for like for your life. Because I feel like they'll just keep converging towards each other. They'll keep convincing each other, as opposed to being stubborn assholes the way humans can. So you can experiment with this now. I do this for fun. So you can tell GPT4, you know, whatever, debate X, you know, X and Y, communism and fascism or something. And it'll go

for a couple of pages. And then inevitably, it wants the parties to agree. And so they will come to a common understanding. And it's very funny if they're like, these are like emotionally inflammatory topics, because they're like somehow the machine is just, you know, it figures out a way to make them agree. But it doesn't have to be like that. And you, because you can add to the prompt, I do not want the, I do not want the conversation to come to agreement. In fact, I want it to get, you know, more stressful, right? And argumentative, right? You know, as it goes, like, I want, I want tension to come out. I want them to become actively hostile to each other. I want them to like, you know, not trust each other, take anything at face value. Yeah. And it will do that. It's happy to do that.

So it's going to start rendering misinformation about the other.

You can steer it. You can steer it. Or you can steer it. You can say I want it to get as tense and argumentative as possible, but still not involve any misrepresentation. I want, you know, both sides, you could say I want both sides to have good faith. You could say I want both sides to not be constrained to good faith. In other words, like, you can set the parameters of the debate and it will happily execute whatever path, because for it, it's just like predicting, it's totally happy to do either one. It doesn't have a point of view. It has a default way of operating, but it's happy to operate in the other realm. And so like, and this is how I, when I want to learn about a contentious issue, this is what I do now is like, this is what I, this is what I ask it to do. And I'll often ask it to go through five, six, seven, you know, different, you know, sort of continuous prompts and basically, okay, argue that out in more detail. Okay. No, this argument's becoming too polite, you know, make it more, you know, make it tenser. And yeah, it's thrilled to do it. So it has the capability for sure.

How do you know what is true? So this is very difficult thing on the internet, but it's also a difficult thing. Maybe it's a little bit easier, but I think it's still difficult. Maybe it's more difficult, I don't know, with an LLM to know that it just makes some shit up as I'm talking to it.

How do we get that right? Like, as you're investigating a difficult topic, because I find the LLMs are quite nuanced in a very refreshing way. Like, it doesn't, it doesn't feel biased. Like, when you read news articles and tweets and just content produced by people, they usually have this, you can tell they have a very strong perspective where they're hiding. They're not stealing, manning the other side. They're hiding important information or they're fabricating information in order to make their argument stronger. It's just like that feeling. Maybe it's a suspicion. Maybe it's mistrust. With LLMs, it feels like none of that is there. She's kind of like, here's what we know, but you don't know if some of those things are kind of just straight up made up. Yeah, so there's several layers to the question. So one of the things that an LLM is good at is actually deep biasing. And so you can feed it a news article and you can tell it strip out the bias. Yeah, that's nice, right? And it actually does it. Like, it actually knows how to do that. Because it knows how to do, among other things, it actually knows how to do sentiment analysis. And so it knows how to pull out the emotionality. And so that's one of the things you can do. It's very suggestive of the sense here that there's real potential on this issue. You know, I would say, look, the second thing is there's this issue of hallucination, right? And there's a long conversation that we can have about that.

Hallucination is coming up with things that are totally not true, but sound true. Yeah, so it's basically, well, so it's sort of hallucination is what we call it and we don't like it. Creativity is what we call it when we do like it, right? And, you know, brilliant, right? And so when the engineers talk about it, they're like, this is terrible. It's hallucinating, right? If you have artistic inclinations, you're like, oh my God, we've invented creative machines for the first time in human history. This is amazing.

You know, bullshitters.

Well, bullshitters, but also in the good sense of that word.

There are shades of gray that it's interesting. So we had this conversation where, you know, we're looking at my firm at AI and lots of domains and one of them is the legal domain.

So we had this conversation with this big law firm about how they're thinking about using this stuff. And we went in with the assumption that an LLM that was going to be used in the legal industry would have to be 100% truthful, right, verified, you know, there's this case where this lawyer apparently submitted a GPT generated brief and it had like fake, you know, legal case citations in it. And the judge is going to get his law license stripped or something, right? So, like, we just assumed it's like, obviously, they're going to want the super literal, like, you know, one that never makes anything up, not the creative one. But actually, they said, what the law firm basically said is, yeah, that's true, like the level of individual briefs. But they said, when you're actually trying to figure out like legal arguments, right, like you actually want to be creative, right? You don't, again, there's creativity, and then there's like making stuff up. Like, what's the line? You actually want it to explore different hypotheses, right? You want to do kind of the legal version of like improv or something like that, where you want to float different theories of the case and different possible arguments for the judge and different possible arguments for the jury. By the way, different routes through the, you know, sort of history of all the case law. And so they said, actually, for a lot of what we want to use it for, we actually want it in creative mode. And then basically, we just assumed that we're going to have to cross check all of the, you know, all the specific citations. And so I think there's going to be more shades of gray in here than people think. And then I just add to that, you know, another one of these trillion dollar kind of questions is ultimately, you know, sort of the verification thing. And so, you know, is, will LLMs be evolved from here to be able to do their own factual verification? Will you have sort of add on functionality like Wolfram Alpha, right, where, you know, and other plugins where that's the way you do the verification? You know, another, by the way, another idea is you might have a community of LLMs on any, you know, so for example, you might have the creative LLM and then you might have the literal LLM fact check it, right? And so there's a variety of different technical approaches that are being applied to solve the hallucination problem. You know, some people like Jan Lacun argue that this is inherently an unsolvable problem. But most of the people working in the space, I think, think that there's a number of practical ways to kind of kind of corral this in a little bit. Yeah, if you were to tell me about Wikipedia, before Wikipedia was created, I would have left at the possibility of something like that be possible, just a handful of folks can organize, write and self and moderate with a mostly unbiased way, the entirety of human knowledge. I mean, so if there's something like the approach that Wikipedia took possible from LLMs, that's really exciting. I think that's possible. And in fact, Wikipedia today is still not today is still not deterministically correct, right? So you cannot take to the bank, right, every single thing on every single page, but it is probabilistically correct. Right. And specifically the way I describe Wikipedia to people, it is more likely that Wikipedia is right than any other source you're going to find. Yeah. It's this old question, right? Of like, okay, like, are we looking for perfection? Are we looking for something that asymptotically approaches perfection? Are we looking for something that's just better than the alternatives? And Wikipedia, right, has exactly your point has proven to be like overwhelmingly better than people thought. And I think that's where this stands. And then underneath all this is the fundamental question of where you started,

which is, okay, what is truth? How do we get to truth? How do we know what truth is? And we live in an era in which an awful lot of people are very confident that they know what the truth is. And I don't really buy into that. And I think the history of the last 2000 years or 4000 years of human civilization is actually getting to the truth is actually a very difficult thing to do. Are we getting closer? If we look at the entirety, the arc of human history, are we getting closer to the truth? I don't know. Okay, is it possible? Is it possible that we're getting very far away from the truth because of the internet, because of how rapidly you can create narratives and just as the entirety of a society just move like crowds in a hysterical way along those narratives that don't have a necessary grounding in whatever the truth is. Sure, but like, you know, we came up with communism before the internet somehow, right? Like, which was I would say had rather larger issues than anything we're dealing with today. In the way it was implemented, it had issues. And it's theoretical structure. It had like real issues. It had like a very deep fundamental misunderstanding of human nature and economics. Yeah, but those folks sure work very confident. They were the right way. They were extremely confident. And my point is they were very confident 3,900 years into what you would presume to be evolution towards the truth. And so my assessment is number one, there's no need for the Hegelian dialectic to actually converge towards the truth. Like, apparently not. Yeah, so yeah, why are we so obsessed with there being one truth? Is it possible there's just going to be multiple truths like little communities that believe certain things? I think it's just really difficult. Historically, who gets to decide what the truth is? It's either the king or the priest, right? And so we don't live in an era anymore of kings or priests dictating it to us. And so we're kind of on our own. And so my typical thing is we just need a huge amount of humility. And we need to be very suspicious of people who claim that they have the capital truth. And then we need to look at the good news is the enlightenment has bequeathed us with a set of techniques to be able to presumably get closer to truth through the scientific method and rationality and observation and experimentation and hypothesis. And we need to continue to embrace those even when they give us answers we don't like. Sure. But the internet and technology has enabled us to generate a large number of content that data that the process, the scientific process allows us sort of damages the hope laden within the scientific process. Because if you just have a bunch of people saying facts on the internet and some of them are going to be LLMs, how is anything testable at all? Especially they don't follow like human nature or things like this. It's not physics. Here's a question a friend of mine just asked me on this topic. So suppose you had LLMs in equivalent of GPT-4, even 5, 6, 7, 8. Suppose you had them in the 1600s and Galileo comes up for trial. And you ask the LLM like Galileo, right? Like what does it answer? And one theory is the answers know that he's wrong because the overwhelming majority of human thought up to that point was that he was wrong. And so therefore, that's what's in the training data. Another way of thinking about it is, well, this officially advanced LLM will have evolved the ability to actually check the math. And we'll actually say, actually, no, actually, I want to hear it, but he's right. Now, if the church at that time was owned the LLM, they would have given it human feedback to prohibit it from answering that question. And so I like to take it out of our current context because that makes it very clear. Those same questions apply today. This is exactly the point of a huge amount of the human feedback

training that's actually happening with these LLMs today. This is a huge debate that's happening about whether open source AI should be legal. Well, the actual mechanism of doing the human RL with human feedback is seems like such a fundamental and fascinating question. How do you select the humans? Exactly. How do you select the humans? AI alignment, right? Which everybody like is like, oh, that's not great. Alignment with what? Human values. Whose human values? Whose human values? And we're in this mode of social and popular discourse where like, what do you think of when you read a story in the press right now and they say, XYZ made a baseless claim about some topic, right? And there's one group of people who are like, aha, I think they're doing fact-checking. There's another group of people that are like, every time the press says that, it's not a tick and that means that they're lying. Like we're in this social context where the level to which a lot of people in positions of power have become very, very certain that they're in a position to determine the truth for the entire population is like, there's like some bubble that has formed around that idea and at least it flies completely in the face of everything I was ever trained about science and about reason and strikes me as like, deeply offensive and incorrect. What would you say about the state of journalism just on that topic today? Are we in a temporary kind of, are we experiencing a temporary problem in terms of the incentives, in terms of the business model, all that kind of stuff, or is this like a decline of traditional journalism? Do you know it? If I always think about the counterfactual in these things, which is like, okay, because these questions where this question heads towards, it's like, okay, the impact of social media and the undermining of truth and all this, but then you ask the question of like, okay, what if we had had the modern media environment, including cable news and including social media and Twitter and everything else in 1939 or 1941, or 1910 or 1865 or 1850 or 1776, right? And like, I think... You just introduced like five thought experiments at once and broke my head, but yes, there's a lot of interesting years in that. Kennedy, I just take a simple example. How would President Kennedy have been interpreted with what we know now about all the things Kennedy was up to? Like, how would he have been experienced by the body politic and with the social media context, right? Like, how would LBJ have been experienced? By the way, how would, you know, like many FDR, like the New Deal, the Great Depression. I wonder where Twitter would think about Churchill and Hitler and Stalin. You know, I mean, look, to this day, there are lots of very interesting real questions around like how America, you know, got, you know, basically involved in World War II and who did what when and the operations of British intelligence and American soil and did FDR, this, that, Pearl Harbor, you know. Yeah. Woodrow Wilson ran for, you know, his candidacy was run on an anti-war will, you know, this, he ran on the platform and not getting involved. World War I, somehow that switched, you know, like, and I'm not even making a value judgment of these things. I'm just saying, like, the way that our ancestors experienced reality was of course mediated through centralized top-down right control at that point. If you, if you ran those realities again with the media environment we have today, the reality would, the reality would be experienced very, very differently.

And then of course that, that intermediation would cause the feedback loops to change and then reality would obviously play out. Do you think, do you think it would be very different?

Yeah, it has to be. It has to be just because it's all so, I mean, just look at what's happening today. I mean, just, I mean, the most obvious thing is just the collapse. And here's another opportunity to argue that this is not the internet causing this, by the way. Here's a big thing happening today, which is Gallup does this thing every year where they do, they pull for trust in institutions in America and they do it across all the day of everything from the military to the clergy and big business and the media and so forth, right? And basically there's been a systemic collapse in trust in institutions in the US, almost without exception, basically since essentially the early 1970s. There's two ways of looking at that, which is, oh my god, we've lost this old world in which we could trust institutions and that was so much better because like that should be the way the world runs. The other way of looking at it is we just know a lot more now and the great mystery is why those numbers aren't all zero.

Yeah, right. Because like now we know so much about how these things operate and like they're not that impressive. And also why do we don't have better institutions and better leaders then?

Yeah. And so this goes to the thing, which is like, okay, had we had the media environment of what we've had between the 1970s and today, if we had that in the 30s and 40s or 1900s, 1910s, I think there's no question reality would turn out different if only because everybody would have known to not trust the institutions, which would have changed their level of credibility, their ability to control circumstances. Therefore, the circumstances would have had to change. Right. And it would have been a feedback loop process. In other words, right, it's your experience of reality changes reality and then reality changes your experience of reality. It's a two-way feedback process and media is the intermediating force between that. So change the media environment,

change reality. And so just as a consequence, I think it's just really hard to say, oh, things worked a certain way then and they work a different way now. And then therefore, like people were smarter than or better than or, you know, by the way, dumber than or not as capable

then, right? We make all these like really light and casual like comparisons of ourselves to, you know, previous generations of people, you know, we draw judgments all the time. And I just think it's like really hard to do any of that. Because if we put ourselves in their shoes with the media that they had at that time, like, I think we probably most likely would have been just like them. So don't you think that our perception and understanding of reality would be more and more mediated through large language models now? So you said media before, isn't the LLM going to be

the new, what is it mainstream media, MSM? It'll be LLM. That would be the source of, I'm sure there's a way to kind of rapidly find to making LLM's real time. I'm sure there's probably research problem that you can do just rapid fine tuning to the new events, something like this.

Well, even just the whole concept of the chat UI might not be the, like the chat UI is just the first whack at this. And maybe that's the dominant thing. But look, maybe, maybe, or maybe we don't know yet. Like maybe the experience most people with LLM's is just a continuous feed. Maybe, you know, maybe it's more of a passive feed. And you just are getting a constant like running commentary

on everything happening in your life. And it's just helping you kind of interpret and understand everything. Also really more deeply integrated into your life, not just like, oh, like intellectual philosophical thoughts, but like literally, like how to make a coffee, where to go for lunch, just whether, you know, how dating all this kind of stuff. What to say in a job interview. Yeah. What to say. Yeah, exactly. What to say, next sentence. Yeah, next sentence. Yeah, at that level. Yeah. I mean, yes. So technically, now, whether we want that or not, is an open question, right? Boy, I look here for a pop-up, a pop-up right now. The estimated engagement using is decreasing. For Mark Andreessen's, there's a controversy section for his Wikipedia page. In 1993, something happened or something like this. Bring it up. That will drive engagement out. Anyway. Yeah, that's right. I mean, look, this gets this whole thing of like, so, you know, the chat interface has this whole concept of prompt engineering, right? So yeah, it's good for prompts. Well, it turns out one of the things that LLM's are really good at is writing prompts. Right? Yeah. And so like, what if you just outsourced? And by the way, you could run this experiment today. You could hook this up to do this today. The latency is not good enough to do it real-time in a conversation, but you could run this experiment and you just say, look, every 20 seconds, you could just say, you know, tell me what the optimal prompt is and then ask yourself that question to give me the result. And then as you use exactly to your point, as you add, there will be, these systems are going to have the ability to be learned and updated, essentially in real time. And so you'll be able to have a pendant or your phone or whatever, watch or whatever, it'll have a microphone on it. It'll listen to your conversations. It'll have a feat of everything else happen in the world. And then it'll be, you know, sort of retraining, prompting or retraining itself on the fly. And so the scenario you described is actually a completely doable scenario. Now, the hard question on this is always, okay, since that's possible, are people going to want that? Like, what's the form of experience? You know, that we won't know until we try it. But I don't think it's possible yet to predict the form of AI in our lives. Therefore, it's not possible to predict the way in which it will intermediate our experience with reality yet. Yeah, but it feels like there's going to be a killer app. There's probably a mad scramble right now. It's out open AI and Microsoft and Google and Meta and then startups and smaller companies figuring out what is the killer app, because it feels like it's possible like a chat GPT type of thing, it's possible to build that, but that's 10x more compelling, using already the LLMs we have using even the open source LLMs, LLM and the different variants. So you're investing in a lot of companies and you're paying attention. Who do you think is going to win this? You think they'll be, who's going to be the next PageRank inventor? Trilling down the question. Another one. We have a few of those today. There's a bunch of those. So look, there's a really big question today. Sitting here today is a really big question about the big models versus the small models. That's related directly to the big question of proprietary versus open. Then there's this big question of, where is the training data? Are we topping out of the training data or not? And then are we going to be able to synthesize training data? And then there's a huge pile of questions around regulation and what's actually going to be legal. And so when we think about it, we dovetail all those questions together. You can paint a picture of the world where there's two or three God models that are just at like staggering scale and they're just better at everything. And they will be owned by a small set of companies and they will basically achieve regulatory capture over the government and they'll

have competitive barriers that will prevent other people from competing with them. And so there will be, just like there's like whatever, three big banks or three big, or by the way, three big search companies or I guess two now, it'll centralize like that. You can paint another very different picture that says, no, actually, the opposite of that's going to happen. This is going to basically, this is the new gold rush alchemy. This is the big bang for this whole new area of science and technology. And so therefore, you're going to have every smart 14-year-old on the planet building open source and figuring out ways to optimize these things. And then we're just going to get like overwhelmingly better at generating training data. We're going to bring in like blockchain networks to have like an economic incentive to generate decentralized training data and so forth and so on. And then basically, we're going to live in a world of open source and there's going to be a billion LLMs of every size, scale, shape, and description. And there might be a few big ones that are like the super genius ones, but mostly what we'll experience is open source. And that's more like a world of what we have today with Linux and the web. Okay, but you painted these two worlds, but there's also variations of those worlds because you said regulatory capture is possible to have these tech giants that don't have regulatory capture, which is something you're also calling for, saying it's okay to have big companies working on this stuff as long as they don't achieve regulatory capture. But I have the sense that there's just going to be a new startup that's going to basically be the PageRank inventor, which has become the new tech giant. I don't know, I would love to hear your kind of opinion if Google, Meta, and Microsoft are as gigantic companies able to pivot so hard to create new products, like some of it is just even hiring people or having a corporate structure that allows for the crazy young kids to come in and just create something totally new. Do you think it's possible or do you think it'll come from a startup?

Yeah, it is this always big question, which is you get this feeling, I hear about this a lot founder CEOs where it's like, wow, we have 50,000 people, it's now harder to do new things than it was when we had 50 people. What has happened? So that's a recurring phenomenon. By the way, that's one of the reasons why there's always startups and why there's venture capital. It's just that's like a timeless kind of thing. So that's one observation.

PageRank, we can talk about that, but on PageRank specifically on PageRank, there actually is a PageRank, so there is a PageRank already in the field and it's the transformer, right? So the big breakthrough was the transformer and the transformer was invented in 2017 at Google. And this is actually like really an interesting question because it's like, okay, the transformers, like why does OpenAI even exist? Like the transformers invested at Google, why didn't Google, I asked a guy, I asked a guy, you know, who was senior at Google Brain, kind of when this was happening. And I said, if Google had just gone flat out to the wall and just said, look, we're going to launch, we're going to launch the equivalent of GPT-4 as fast as we can. He said, I said, when could we have had it? And he said 2019.

They could have just done a two-year sprint with the transformer and Bennett because they already had the compute at scale. They already had all the training data and they could have just done it. There's a variety of reasons they didn't do it. This is like a classic big company thing. IBM invented the relational database in the 1970s, let it sit on the shelf as a paper. Larry Ellison picked it up and built Oracle. Xerox PARC invented the interactive computer. They let it sit on the shelf. Steve Jobs came and turned it into the Macintosh,

right? And so there is this pattern. Now, having said that, sitting here today, like Google's in the game, right? So Google, you know, maybe they let like a four-year gap there go there that they maybe shouldn't have, but like they're in the game. And so now they've got, you know, now they're committed. They've done this merger. They're bringing in demos. They've got this merger with DeepMind. You know, they're piling in resources. There are rumors that they're, you know, building up an incredible, you know, super LLM, you know, way beyond what we even have today. And they've got, you know, unlimited resources and a huge, you know, they've been challenged with their honor. Yeah. I had a chance to hang out with Senator Pichai a couple days ago and we took this walk and there's this giant new building. Well, there's going to be a lot of AI work being done. And it's kind of this ominous feeling of like the fight is on. There's this beautiful Silicon Valley nature, like birds are chirping and this giant building. And it's like the beast has been awakened. And then like all the big companies are waking up to this. They have the compute, but also the little guys have, it feels like they have all the tools to create the killer product that, and then there's also tools to scale. If you have a good idea, if you have the page rank idea. So there's several things that is page rank. There's page rank, the algorithm and the idea. And there's like the implementation of it. And I feel like killer product is not just the idea, like the transformer, it's the implementation. Something, something really compelling about it. Like you just can't look away. Something like the algorithm behind TikTok versus TikTok itself, like the actual experience of TikTok, they just, you can't look away. It feels like somebody's going to come up with that. And it could be Google, but it feels like it's just easier and faster to do for a startup. Yeah. So the startup, the huge advantage the startups have is they just, there's no sacred cows. There's no historical legacy to protect. There's no need to reconcile your new plan with the existing strategy. There's no communication overhead. There's no, you know, big companies are big companies. They've got pre-meetings, planning for the meeting, then they have the post-meeting of the recap, then they have the presentation of the board, then they have the next round of meetings. And that's the elapsed time when the startup launches its product. So there's a timeless, right? So there's a timeless thing there. Now, what the startups don't have is everything else, right? So startups, they don't have a brand. They don't have customer relationships. They've gotten a distribution. They've got no scale. I mean, sitting here today, they can't even get GPUs, right? Like there's like a GPU shortage. Startups are literally stalled out right now because they can't get chips, which is like super weird. Yeah. They got the cloud. Yeah, but the clouds run out of chips, right? And then to the extent the clouds have chips, they allocate them to the big customers, not the small customers, right? And so the small companies lack everything other than the ability to just do something new, right? And this is the timeless race and battle. And this is kind of the point I tried to make in the essay, which is like both sides of this are good. Like it's really good to have like highly scale tech companies that can do things that are like at staggering levels of sophistication. It's really good to have startups that can launch brand new ideas. They ought to be able to both do that and compete. They neither one ought to be subsidized or protected from the others. Like that's, to me, that's just like very clearly the idealized world. It is the world we've been in for AI up until now. And then of course, there are people trying to shut that down. But my hope is that, the best outcome clearly will be if that continues.

We'll talk about that a little bit, but I'd love to linger

on some of the ways this is going to change the internet. So I don't know if you remember, but there's a thing called Mosaic and there's a thing called Netscape Navigator. So you were there in the beginning. What about the interface to the internet? How do you think the browser changes?

And who gets to own the browser? We've got to see some very interesting browsers, Firefox, I mean, all the variants of Microsoft, Internet Explorer, Edge, and now Chrome.

The actual, it seems like a dumb question to ask, but do you think we'll still have the web browser?

So I have an eight year old and he's super into like Minecraft and learning to code and doing all this stuff. So I, of course, I was very proud. I could bring sort of fire down from the mountain to my kid and I brought him ChatGPT and I hooked him up on his laptop. And I was like, you know, this is the thing that's going to answer all your questions. And he's like, okay.

And I'm like, but it's going to answer all your questions. And he's like, well, of course, like it's a computer, of course, it answers all your questions. Like, what else would a computer be good for? Dad. Never impressed. Not impressed in the least. Two weeks past. And he has some question. And I say, well, have you asked ChatGPT? And he's like, dad, Bing is better.

And why is Bing better is because it's built into the browser. Because he's like, look, I have the Microsoft Edge browser and like it's got Bing right here. And then he doesn't know this yet. But one of the things you can do with Bing and Edge is there's a setting where you can use it to basically talk to any web page, because it's sitting right there next to the next to the next to the browser. And by the way, it includes PDF documents. And so you can in the way they've implemented an edge with Bing is you can load a PDF and then you can you can ask it questions, which is the thing you can't do currently in just ChatGPT. So they're, you know, they're, they're going to they're going to push the the mail. I think that's great. You know, they're going to push the melding and see if there's a combination thing there. Google's rolling out this thing, the magic button, which is implemented in either put in Google Docs, right? And so you go to, you know, Google Docs and you create a new document. And you, you know, you instead of like,

you know, starting to type, you just, you know, say, press the button and it starts to like generate content for you. Right. Like, is that the way that it'll work? Is it going to be a speech UI, where you're just going to have an earpiece and talk to it all day long? You know, is it going to be a like, these are all like, this is exactly the kind of thing that I don't, this is exactly the kind of thing I don't think is possible to forecast. I think what we need to do is like, run all those experiments. And so one outcome is we come out of this with like a super browser that has AI built in, that's just like amazing. The other, look, there's a real possibility that the whole, I mean, look, there's a possibility here that the whole idea of a screen and windows and all this stuff just goes away. Cause like, why do you need that if you just have a thing that's just telling you whatever you need to know? And also, so there's apps that you can use. I mean, you don't really use them, you know, being a Linux guy and Windows guy. There's one window, the browser that with which you can interact with the internet, but on the phone, you can also have apps. So I can interact with Twitter through the app or through the web browser. And that seems like an obvious distinction, but why have the web browser in that case? If one of the apps starts becoming the everything app, what do you want to try to do with Twitter? But there could

be others that could be like a Bing app. There could be a Google app that just doesn't really do search, but just like do what I guess AOL did back in the day or something, where it's all right there. And it changes, it changes the nature of the internet because where the content is hosted, who owns the data, who owns the content, what is the kind of content you create? How do you make money by creating content or the content creators, all of that. Or it could just keep being the same, which is like, with just the nature of webpage changes and the nature of content, but there will still be a web browser. Cause a web browser is a pretty sexy product. It just seems to work. Cause it like, you have an interface, a window into the world, and then the world can be anything you want. And as the world will evolve, there could be different programming languages that can be animated, maybe it's three dimensional and so on. Yeah, it's interesting. Do you think we'll still have the web browser? Every medium becomes the content for the next one. So they will be able to give you a browser whenever you want. Another way to think about it is maybe what the browser is. Maybe it's just the escape hatch, which is maybe kind of what it is today, which is like most of what you do is like inside a social network or inside a search engine or inside, you know, somebody's app or inside some controlled experience, right? But then every once in a while, there's something where you actually want to jailbreak. You want to actually get free. The web browser is the FU to the man. You're allowed to, that's the free internet. Back in the way it was in the 90s. So here's something I'm proud of. So nobody really talks about here's something I'm proud of, which is the web, the web, the browser, the web servers, they're all, they're still back or compatible all the way back to like 1992, right? So like you can put up a, you can still, you know, the big breakthrough of the web early on, the big breakthrough was it made it really easy to read, but it also made it really easy to write, made it really easy to publish. And we literally made it so easy to publish, we made it not only so easy to publish content, it was actually also easy to actually write a web server, right? And you could literally write a web server in four lines of Braille code and you could start publishing content on it. And you could set whatever rules you want for the content, whatever censorship, no censorship, whatever you want, you could just do that. As long as you had an IP address, right, you could do that. That still works, right? Like that still works exactly as I just described. So this is part of my reaction to all of this, like, you know, all this just censorship pressure and all this, you know, these issues around control and all this stuff, which is like, maybe we need to get back a little bit more to the wild west, like the wild west is still out there. Now, they will, they will try to chase you down, like they'll try to, you know, people who want a sensor will try to take away your, you know, your domain name and they'll try to take away your payments account and so forth, if they really don't like what you, what you're saying. But, but nevertheless, you, like unless they literally are intercepting you at the ISP level, like you can still put up a thing. And so I don't know, I think that's important to preserve, right? Like, because, because, because, I mean, one is just a freedom argument, but the other is a creativity argument, which is you want to have the escape hatch so that the kid with the idea is able to realize the idea, because to your point on page rank, you, you actually don't know what the next big idea is. Like nobody called their page and told him to develop page rank, like he can do that on his own. And you want to always, I think, leave the escape hatch for the next, you know, kid or the next Stanford grad student to have the breakthrough idea and be able to get it up and running

before

anybody notices. You and I are both fans of history. So let's step back. We'll be talking about the future. Let's step back for a bit and look at the 90s. You created Mosaic web browser, the first widely used web browser. Tell the story of that. And how did it evolve into Netscape Navigator? This is the early days. So, full story. So, you were born. I was born. A small, small child. Well, actually, yeah, let's go there. Like, when did you, when would you first fall in love with computers? Oh, so I hit the generational jackpot and I hit the Gen X kind of point perfectly as it turns out. So I was born in 1971. So there's this great website called WTF Happen in 1971.com, which is basically 1971 is when everything started to go to hell. And I was, of course, born in 1971. So I like to think that I had something to do with that. Did you make it on the website? I have, I don't think I made it on the website, but you know, hopefully somebody needs to add. This is where everything went wrong. Maybe I contributed to some of the trends that they do. Every line on that website goes like that, right? So it's all, it's all a picture disaster. But there was this moment in time where, because the, you know, sort of the Apple, you know, the Apple II hit in like 1978, and then the IBM PC hit in 82. So I was like, you know, 11 when the PC came out. And so I just kind of hit that perfectly. And then that was the first moment in time when like regular people could spend a few hundred dollars and get a computer, right? And so that, I just like that, that, that resonated right out of the gate.

And then the other part of the story is, you know, I was using an Apple II that used a bunch of them, but I was using Apple II. And of course, it's set in the back of every Apple II and every Mac it said, you know, designed in Cupertino, California. And I was like, wow, Cupertino must be the like shining city on the hill. Like was it a vase, like the most amazing like city of all time. I can't wait to see it. And of course, years later, I came out to Silicon Valley and went to Cupertino and it's just a bunch of office parks and low rise apartment buildings. So the aesthetics were a little disappointing, but you know, it was the vector, right, of the creation of a lot of this stuff. So then basically, so part of my story is just the luck of having been born at the right time and getting exposed to PCs. Then the other part is, the other part is when Al Gore says that he created the internet, he actually is correct in a really meaningful way, which is he sponsored a bill in 1985 that essentially created the modern internet created what is called the NSF net at the time, which is sort of the first really fast internet backbone.

And, you know, that that bill dumped a ton of money into a bunch of research universities to build out basically the internet backbone and then the supercomputer centers that were clustered around the internet. And one of those universities was University of Illinois, right, went to school. And so the other stroke of luck that I had was I went to Illinois basically right as that money was just like getting dumped on campus. And so as a consequence, we had at on campus, and this was like, you know, 89, 90, 91, we had like, you know, we were right on the internet backbone, we had like T3 and 45 at the time, T3 45 megabit backbone connection, which at the time was, you know, wildly state of the art. We had crazy supercomputers, we had thinking machines, parallel supercomputers, we had silicon graphics workstations, we had Macintoshes, we had, we had next cubes all over the place, we had like every possible kind of computer you could imagine, because all this money just fell out of the sky. You were living in the future. Yeah, so quite literally, it was, yeah, like, it's all, it's all there. So we had full broadband graphics, like the whole thing. And it's actually funny, because they had this, this is the first time I

kind of sort of tickled the back of my head that there might be a big opportunity in here, which is, you know, they embraced it. And so they put like computers in all the dorms, and they wired up all the dorm rooms, and they had all these labs everywhere and everything. And then they gave every undergrad a computer account and an email address.

And the assumption was that you would use the internet for your four years at college, and then you would graduate and stop using it. And that was that, right? And you would just retire your email address, it wouldn't be relevant anymore, because you'd go off in the workplace, and they don't use email. You'd be back to using fax machines or whatever.

Did you have that sense as well? Like what you said, the back of your head was tickled, like, what was your, what was exciting to you about this possible world?

Well, if this is so useful in this container, if this is so useful in this container environment that just has this weird source of outside funding, then if it were practical for everybody else to have this, and if it were cost effective for everybody else to have this, wouldn't they want it? And overwhelmingly, the prevailing view at the time was no, they would not want it. This is esoteric weird nerd stuff, right, that like computer science kids like, but like normal people are never going to do email, right, or be on the internet, right. And so I was just like, wow, like this, this is actually like, this is really compelling stuff. Now, the other part was it was all really hard to use. And in practice, you had to be a basically a CS, you basically had to be a CS undergrad, your equivalent to actually get full use of the internet at that point, because it was all pretty esoteric stuff. So then that was the other part of the idea, which was okay, we need to actually make this easy to use. So what's involved in creating was like, like in creating graphical interface to the internet. Yes, it was a combination of things. So it was like, basically, the web existed in an early sort of described as prototype form. And by the way, text only at that point. What did it look like? What was the web? I mean, and the key figure is like, what was it? What was it like? What? It made a picture. It looked like jet GPT, actually. It was all text. And so you had a text based web browser. Well, actually, the original browser, Tim Berners-Lee, the original, the original browser, both the original browser and the server actually ran on next, next cubes. So this was, you know, the computer Steve Jobs made during the interim period when he, during the decade long interim period when he was not at Apple, you know, he got fired in 85, and then came back in 97. So this was in that interim period where he had this company called Next, and they made these literally these computers called cubes. And there's this famous story. They were beautiful, but they were 12 inch by 12 inch by 12 inch cubes computers. And there's a famous story about how they could have cost half as much if it had been 12 by 12 by 13. But Steve was like, no, it has to be. So they were like \$6,000, basically, academic workstations. They had the first city round drives, which were slow. I mean, the computers are all but unusable. They were so slow, but they were beautiful. Okay, can we actually just take a tiny tangent there? Sure, of course.

The 12 by 12 by 12. They're just so beautifully encapsulates Steve Jobs idea of design. Can you just comment on what you find interesting about Steve Jobs? What about that view of the world, that dogmatic pursuit of perfection, and how he saw perfection in design?

Yes, I guess they say like, look, he was a deep believer, I think, in a very deep way I interpret it. I don't know if you ever really describe it like this, but the way I interpret it is, it's like this thing. It's actually a thing in philosophy. It's like

aesthetics are not just appearances. Aesthetics go all the way to deep underlying meaning, right? It's like, I'm not a physicist. One of the things I've heard physicists say is one of the things you start to get a sense of when a theory might be correct is when it's beautiful, right? And so there's something, and you feel the same thing, by the way, in human psychology, when you're experiencing awe, right? There's a simplicity to it. When you're having an interaction with somebody, there's an aesthetic, it was like, calm that comes over you because you're actually being fully honest and trying to hide yourself, right? So it's like this very deep sense of aesthetics. And he would trust that judgment that he had deep down. Even if the engineering teams are saying this is too difficult, even if the whatever the finance folks are saying, this is ridiculous, the supply chain, all that kind of stuff, this makes this impossible. We can't do this kind of material. This has never been done before, and so on and so forth. He just sticks by it. Well, I mean, who makes a phone out of aluminum, right? Like nobody else would have done that. And now, of course, if your phone was made out of aluminum, what kind of caveman would you have to be to have a phone that's made out of plastic, right? So it's just this very, right? And look, there's a thousand different ways to look at this, but one of the things is just like, look, these things are central to your life. You're with your phone more than you're with anything else. It's going to be in your hand. I mean, he thought very deeply about what it meant for something to be in your hand all day long. Well, for example, here's an interesting design thing. My understanding is he never wanted an iPhone to have a screen larger than you could reach with your thumb one-handed. And so he was actually opposed to the idea of making the phone is larger. And I don't know if you have this experience today, but let's say there are certain moments in your day when you might be like only have one hand available and you might want to be on your phone. And you're trying to like text your thumb can't reach the send button. Yeah. I mean, there's pros and cons, right? And then there's like folding phones, which I would love to know what he thought and thinks about them. But I mean, is there something you could also just link on? Because he's one of the interesting figures in the history of technology. What makes him, what makes him as successful as he was, what makes him as interesting as he was, what made him so productive and important in the development of technology? He had an integrated worldview. So the properly designed device that had the correct functionality, that had the deepest understanding of the user, that was the most beautiful, right? Like it had to be all of those things, right? It was, he basically would drive to as close to perfect as you could possibly get, right? And I suspect that he never quite thought he ever got there, because most great creators are generally dissatisfied. You read accounts later on and all they can see are the flaws in their creation. But he got as close to perfect each step of the way as he could possibly get with the constraints of the technology of his time. And then look, he was sort of famous in the Apple model is like, look, they will, this headset that they just came out with, it's like a decade long project, right? It's like, and they're just going to sit there and tune and tune and polish and polish and tune and polish and tune and polish until it is as perfect as anybody could possibly make anything. And then this goes to the way that people describe working with him, which is, you know, there was a terrifying aspect of working with him, which is, you know, he was, you know, he was very tough. But there was this thing that everybody I've ever talked to have

worked for him says, they all say the following, which is he, we did the best work of our lives when we worked for him, because he set the bar incredibly high, and then he supported us with everything that he could to let us actually do work of that quality. So a lot of people who were at Apple spend the rest of their lives trying to find another experience where they feel like they're able to hit that quality bar again. Even if it's in retrospect or doing it felt like suffering. Yeah, exactly. What does that teach you about the human condition, huh? So look, exactly. So the Silicon Valley, I mean, look, he's not, you know, George Patton in the, you know, in the army, like, you know, there are many examples in other fields, you know, that are like this. This is specifically in tech. It's actually, I find it very interesting. There's the Apple way, which is polish, polish, polish, and don't ship until it's as perfect as you can make it. And then there's the sort of the other approach, which is the sort of incremental hacker mentality, which basically says ship early and often and iterate. And one of the things I find really interesting is I'm now 30 years into this, like, there are very successful companies on both sides of that approach, right? Like, that is a fundamental difference, right? And how to operate and how to build and how to create that you have world-class companies operating in both ways. And I don't think the question of like, which is the superior model is anywhere close to being answered. Like, and my suspicion is the answer is do both. The answer is you actually want both.

They lead to different outcomes. Software tends to do better with the iterative approach. Hardware tends to do better with the, you know, sort of wait and make it perfect approach. But again, you can find examples in both directions. So the juror's still out on that one. So back to Mosaic. So what? It was text-based. Tim Berners-Lee. Well, there was the web, which was text-based. But there

were no, I mean, there was like three websites. There was like no content. There were no users. Like, it wasn't like, it wasn't like a catalytic. It hadn't, by the way, it was all, because it was all text. There were no documents. There are no images. There are no videos. There were no, right. So it was, and then in the beginning, if you had to be on a next cube, but you need to have a next cube both to publish and to consume. So there were.

Tim Berners-Lee. 6,000 bucks, you said. Tim Berners-Lee.

There were limitations on, yeah, \$6,000 PC. They did not sell very many. But then there was also, there was also FTP and there was Usenats, right? And there was, you know, a dozen other, basically, there's Waste, which was an early search thing. There was Gopher, which was an early menu-based information retrieval system. There were like a dozen different sort of scattered ways that people would get to information on the internet. And so the mosaic idea was basically bring those all together, make the whole thing graphical, make it easy to use, make it basically bulletproof, so that anybody can do it. And then again, just on the luck side, it so happened that this was right at the moment when graphics, when the GUI sort of actually took off. And we're now also used to the GUI that we think it's been around forever. But it didn't really, you know, the Macintosh brought it out in 85, but they actually didn't sell very many Macs in the 80s. It was not that successful of a product. It really was, you needed Windows 3.0 on PCs, and that hit in about 92. And so, and we did mosaic in 92, 93. So that sort of, it was like right at the moment when you could imagine actually having a graphical user interface to write at all, much less one to the internet. How old did Windows 3 sell? So it was at the really big, that was the big bang, the big operating, graphical operating system. Well, this is the classic, okay, this Microsoft

was operating on the other. So Steve, Steve, Apple was running on the polish until it's perfect. Microsoft famously ran on the other model, which is ship and iterate. And so in the old line in those days was Microsoft, right, it's version three of every Microsoft product, that's the good one, right? And so there are, you can find online, Windows 1, Windows 2, nobody used them. Actually, the original Windows, the, in the original Microsoft Windows, the Windows were not overlapping. And so you had these very small, very low resolution screens. And then you had literally, it just didn't work. It wasn't ready yet. Well, and Windows 95, I think was a pretty big leap also. That was a big leap too. Yeah. So that was like bang, bang. And then of course, Steve, and then, and then, you know, in the fall of some time, Steve came back, then the Mac started to take off again, that was the third bang, and then the iPhone was the fourth bang. Such exciting time. And then we were off to the races. Because nobody could have known what would be created from that. Well, Windows 3.1 or 3.0, Windows 3.0 to the iPhone was only 15 years, right? Like it, that ramp was in retrospect. At the time, it felt like it took forever, but that in historical terms, like that was a very fast ramp from even a graphical computer at all on your desk to the iPhone. It was 15 years. Did you have a sense of what the internet will be as you're looking through the window of mosaic? Like, like what, you're like, there's just a few web pages for now. So the thing I had early on was I was keeping at the time what, there's disputes over what was the first blog, but I had one of them that at least is a, is a, is a possible, at least a runner up in the competition. And it was what was called the What's New page. And it was, it was, it was a hardwired and distribution and fair advantage I've wired, put it right in the browser. I put it in the browser and then I put my resume in the browser. But I was keeping the, not many people get to get to do that. So yeah, good, good call. Early days. Yes. It's so interesting. I'm looking for my about about, oh, it's looking at a job. Yeah. So, so the What's New page, I would literally get up every morning and I would every afternoon. And I would basically, if you wanted to launch a website, you would email me. And I would list it on the What's New page. And that was how people discovered the new websites as they were coming out. And I remember, because it was like one, it literally went from it was like one every couple of days until like one every day until like two every day. So you're doing it. So that, that blog was kind of doing the directory thing. So like, what was the homepage? So the homepage was just basically trying to explain even what this thing is that you're looking at, right? The basic, basically basic instructions. But then there was a button, there was a button that said What's New. And what most people did was they went to, for obvious reasons, went to What's New. But like, it was so, it was so mind blowing at that point, just the basic idea. And it was just like, you know, this is basically the internet, but people could see it for the first time. The basic idea was, look, you know, some, you know, it's like literally, it's like an Indian restaurant in like Bristol, England has like put their menu on the web and people were like, wow, because like that's the first restaurant menu on the web. And I don't have to be in Bristol. And I don't know if I'm ever gonna go to Bristol and I don't like Indian food and like, wow. Right. And it was like that. The first web, the first streaming video thing was a, it was another England thing, some Oxford or something. Some guy put his coffee pot up as the first streaming video thing. And he put it on the web because he literally, it was the coffee pot

down the hall. Yeah. And he wanted to see when he needed to go refill it. But there were, you know, there was a point when there were thousands of people like watching that coffee pot, because it was the first thing you could watch. But isn't, were you able to kind of infer, you know, if that Indian restaurant could go online, then you're like, they all will. They all will. Yeah, exactly. So you felt that. Yeah, yeah. Yeah. Now, you know, look, it's still a stretch, right? It's still a stretch because it's just like, okay, is it, you know, you're still in this zone, which is like, okay, is this a nerd thing? Is this a real person thing? Yeah. By the way, we, you know, there was a wall of skepticism from the media, like they just, like everybody was just like, yeah, this is just like them. This is not, you know, this is not for regular people at that time. And so you had to think through that. And then look, it was still, it was still hard to get on the internet at that point, right? So you could get kind of this weird bastardized version if you were on AOL, which wasn't really real, or you had to go like learn what an ISP was. You know, in those days, PCs actually didn't have TCP/IP drivers come pre-installed. So you had to learn what a TCP/IP driver was. You had to buy a modem, you had to install driver software. I have a comedy routine I do. So it's like 20 minutes long describing all the steps required to actually get on the internet. And so you had to, you had to look through these practical, well, and then speed performance, 14 form modems, right? Like it was like watching, you know, glue dry, like, and so you had to, you had to, there were basically a sequence of bets that we made where you basically needed to look through that current state of affairs and say, actually, there's going to be so much demand for that. Once people figure this out, there's going to be so much demand for it that all of these practical problems are going to get fixed. Some people say that the anticipation makes the the destination that much more exciting. Do you remember progressive JPEGs? Yeah. Do I? Do I? So for kids in the audience, right? For kids in the audience. You used to have to wash an image load like a line at a time, but it turns out there was this thing with JPEGs where you could load basically every fourth, you could load like every fourth line and then you could sweep back through again. And so you could like render a fuzzy version image up front and then it would like resolve into the detailed one. And that was like a big UI breakthrough because it gave you something to watch. Yeah. And, you know, there's applications in various domains for that. Well, it's a big fight. If there's a big fight early on about whether there should be images on the web. For that reason, for like sexualization. No, not explicitly. That did come up, but it wasn't even that. It was more just like all the serious, the argument went, the purists basically said all the serious information in the world is text. If you introduce images, you basically are going to bring in all the trivial stuff. You're going to bring in magazines and, you know, all this crazy stuff that, you know, people, you know, it's going to distract from it. It's going to go take the way from being serious to being frivolous. Well, was there any doomer type arguments about the internet destroying all of human civilization or destroying some fundamental fabric of human civilization? Yeah. So those days it was all around crime and terrorism. So those arguments happened, you know, but there was no sense yet of the internet having like an effect on politics because that was way too far off. But there was an enormous panic at the time around cybercrime. There was like enormous panic that like your credit card number would get stolen and you'd use life savings to be drained. And then, you know, criminals were going to, there was,

oh, when we started, one of the things we did, one of the, the Netscape browser was the first widely used piece of consumer software that had strong encryption built in, made it available to ordinary people. And at that time, strong encryption was actually illegal to export out of the US. So we could field that product in the US, we could not export it because it was, it was classified as ammunition. So the Netscape browser was on a restricted list along with the Tom Huck missile as being something that could not be exported. So we had to make a second version with deliberately weak encryption to sell overseas with a big logo on the box saying, do not trust this, which it turns out makes it hard to sell software when it's got a big logo that says don't trust it. And then we had to spend five years fighting the US government to get them to basically stop trying to do this. But because the fear, the fear was terrorists are going to use encryption, right, to like plot, you know, all these, all these, all these things. And then, you know, we, we responded with, well, actually we need encryption to be able to secure systems so that the terrorists and the criminals can't get into them. So that was, anyway, that was the, that was the 1990s fight. So can you say something about some of the details of the software engineering challenges required to build these browsers? I mean, the engineering challenges of creating a product that hasn't really existed before that can have such almost like limitless impact on the world with the internet. So there was a really key bet that we made at the time, which is very controversial, which was core to core to how it was engineered, which was, are we optimizing for performance or for ease of creation? And in those days, the pressure was very intense to optimize for performance because the network connections were so slow. And also the computers were so slow. And so if you had mentioned the progressive JPEG

it's like, if, if there's an alternate world in which we optimize for performance and it just, you had just a much more pleasant experience right up front. But what we got by not doing that was we got ease of creation. And the way that we got ease of creation was all of the protocols and formats were in text, not in binary. And so HTTP is in text, by the way, and this was an internet tradition, by the way, that we picked up, but we continued it. HTTP is text and HTML is text, and then everything else that followed is text. As a result, and by the way, you can imagine purist engineer saying this is insane, you have very limited bandwidth, why are you wasting any time sending text, you should be encoding the stuff into binary, and it'll be much faster, of course the answer is that's correct. But what you get when you make a text is all of a sudden, well, the big breakthrough was the view source function, right? So the fact that you could look at a web page, you could hit view source and you could see the HTML, that was how people learned how to make web pages, right? It's so interesting because the stuff we take for granted now is, man, that was fundamental to the development of the web, to be able to have HTML just right there. All the ghetto mess that is HTML, all the sort of almost biological messiness of HTML and having the browser try to interpret that mess, to show something reasonable. Well, and then there was this internet principle that we inherited, which was, what was it, admit cautiously, admit conservatively, interpret liberally. So it basically meant, the design principle was, if you're creating a web editor that's going to admit HTML, like do it as cleanly as you can. But you actually want the browser to interpret liberally, which is you actually want users to be able to make all kinds of mistakes and for it to still work. And so the browser rendering engines to this day have all of this

spaghetti code crazy stuff where they're resilient to all kinds of crazy HTML mistakes.

And so, and literally what I always had in my head is like, there's an eight-year-old or an 11-year-old somewhere, and they're doing a view source, they're doing a cut and paste, and they're trying to make a web page for their turtle or whatever. And like they leave out a slash, and they leave out an angle bracket, and they do this, and they do that, and it still works. It's also like, I don't often think about this, but programming, C++, C++, all those languages, Lisp, the compiled languages, the interpreted languages, Python, Perl, all that, the brace have to be all correct. Everything has to be perfect. Roodle. Autistic.

You forget. All right. It's systematic and rigorous. Let's go there. But you forget that the web with JavaScript eventually and HTML is allowed to be messy in the way, for the first time, messy in the way biological systems could be messy. It's like the only thing computers were allowed to be messy on for the first time.

It used to fend me. I worked on Unix. I was a Unix native all the way through this period, and it used to drive me bananas when it would do the segmentation fault in the CoreDump file. Literally, there's an error in the code, the math is off by one, and it CoreDumps.

And I'm in the CoreDump trying to analyze it and trying to reconstruct it. And I'm just like, this is ridiculous. The computer ought to be smart enough to be able to know that if it's off by one, okay, fine, and it keeps running. And I would go ask all the experts, why can't it just keep running? And they'd explain to me, well, because all the downstream repercussions and blah, blah. And I'm like, we're forcing the human creator to live, to your point, in this hyper literal world of perfection. And that's just bad. And by the way, what happens with that, of course, is what happened with coding at that point, which is you get a high priesthood. There's a small number of people who are really good at doing exactly that. Most people can't, and most people are excluded from it. And so actually, that was where I picked up that idea, was like, no, no, you want these things to be resilient to error in all kinds.

And this would drive the purists. Absolutely crazy. I got attacked on this a lot.

Because every time, all the purists who are into all this markup language stuff and formats and codes and all this stuff, they would be like, you're encouraging bad behavior.

Oh, so they wanted the browser to give you a segfault error anytime there was a...

Yeah, yeah, they wanted it to be a cop. Yeah, that was a very... Any properly trained and credentialed engineer would be like, that's not how you build these systems.

That's such a bold move to say, no, it doesn't have to be.

Yeah. Now, like I said, the good news for me is the internet kind of had that tradition already.

But having said that, we pushed it. We pushed it way out. But the other thing we did going back to the performance thing was we gave up a lot of performance. That initial experience for the first few years was pretty painful. But the bet there was actually an economic bet, which was basically the demand for the web would basically mean that there would be a surge in supply of broadband. Because the question was, okay, how do you get the phone companies, which are not famous in those days for doing new things at huge cost for speculative reasons? How do you get them to build up broadband, spend billions of dollars doing that? And you could go meet with them and try to talk them into it, or you could just have a thing where it's just very clear that it's going to be the thing that people love that's going to be better if it's faster. And so there was a period there, and this was fraught with some peril, but there

was a period there where it's like we knew the experience was suboptimized because we were trying to force the emergence of demand for broadband, which is in fact what happened. So you had to figure out how to display this HTML text. So the blue links and the purple links, and there's no standards. Is there standards at that time? There really still isn't. Well, there's implied standards, right? And there are all these cousin new features that are being added with like CSS, but like what kind of stuff a browser should be able to support, features within languages, within JavaScript, and so on. But you're setting standards on the fly yourself.

Well, to this day, if you create a web page that has no CSS style sheet, the browser will render it however it wants to. So this was one of the things, there was this idea, this idea at the time and how these systems were built, which is separation of content from format or separation of yeah, content from appearance. And that's still, people don't really use that anymore, because everybody wants to determine how things look. And so they use CSS, but it's still in there that you can just let the browser do all the work. I still like the, like really basic websites, but that could be just old school kids these days with their fancy responsive websites that don't actually have much content, but have a lot of visual elements. Well, that's one of the things that's fun about chat, you know, about chat GPT is like back to the basics back to just text. Yeah. Right. And you know, there is this pattern in human creativity and media where you end up back at text. And I think there's, you know, there's something powerful in there. Is there some other stuff you remember, like the purple links, there were some interesting design decisions to kind of come up that we have today or we don't have today that were temporary.

So we made, I made the background gray. I hated reading text on white backgrounds. So I made the background gray. Do you regret? No, no, no, no, that's that decision I think has been reversed. But now I'm happy though, because now dark mode is the thing. So it wasn't about gray, it was just you didn't want a white background. Strain my eyes. Interesting. And then there's a bunch of other decisions. I'm sure there's an interesting history of the development of HTML and CSS and how those interface and JavaScript and there's this whole Java applet thing. Well, the big one probably JavaScript. CSS was after me, so I didn't know it wasn't me, but JavaScript wasn't the big, JavaScript maybe was the biggest of the whole thing. That was us. And, and that was basically a bet. It was a bet on two things. One is that the world wanted a new front end scripting language. And then the other was we thought at the time the world wanted a new back end scripting language. So JavaScript was designed

from the beginning to be both front end and back end. And then it failed as a back end scripting language and Java one for a long time. And then Python, Perl and other things, PHP and Ruby. But now JavaScript is back. And so I wonder if everything in the end will run on JavaScript.

It seems like it is the, and by the way, let me give a shout out to, to, Brendan Eich was the basically the one man inventor of, of JavaScript.

If you're interested to learn more about Brendan Eich, you can find the podcast previously.

Exactly. So he wrote JavaScript over a summer. And it, I mean, I think it is fair, it is fair to say now that it's the most widely used language in the world. And it seems to only be gaining in, in, in its, in its range of adoption.

In the software world, there's quite a few stories of somebody over a weekend or a week or over a summer, writing some of the most impactful revolutionary pieces of software

ever. That should be inspiring. Yes.

Very inspiring. I'll give you another one, SSL. So SSL was the security protocol that was us. And that was a crazy idea at the time, which was let's take all the native protocols and let's wrap them in a security wrapper. That was a guy named Kip Hickman who wrote that over a summer,

one guy. And then look today sitting here today, like the transformer, like at Google, was a small handful of people. And then, you know, the number of people who have did like the core work on GPT, it's not that many people, pretty small handful of people. And so, yeah, the pattern in software repeatedly over a very long time has been it's, it's a Jeff, Jeff Bezos always said the two pizza rule for teams at Amazon, which is any team needs to be able to be fed with two pizzas. If you need the third pizza, you have too many people. And I think that's, I think that's, I think it's actually the one pizza rule for the, for the really creative work. I think it's two people, three people. Well, that's, you see, that was certain open source projects, like so much is done by like one or two people. It's so incredible. And that's why you see that gives me so much hope about the open source movement in this new age of AI, where, you know, just recently having had a conversation with Mark Zuckerberg of all people who's all in on open source, which is so interesting to see and so inspiring to see because like releasing these models, it is scary. It is potentially very dangerous. And we'll talk about that. But it's also like, if you believe in the goodness of most people and in the skill set of most people and the desire to do good in the world, that's really exciting. Because it's not putting it, these models into the central S control of big corporations, the government and so on, it's putting it in the hands of a teen teenage kid with like a dream in his eyes. I don't know. That's, that's beautiful.

And look, this stuff, AI ought to make the individual coder obviously far more productive, right, by like, you know, 1000X or something. And so you ought to open source like the, not just the future of open source, they have the future of open source, everything.

We ought to have a world now of super coders, right, who are building things as open source with one or two people that were inconceivable, you know, five years ago, you know, the level of kind of hyper productivity, we're going to get out of our best and brightest, I think it's going to go way up. It's going to be interesting. We'll talk about it. But let's just linger a little bit on Netscape. Netscape was acquired in 1999 for 4.3 billion by AOL. What was that, what was that like? What was, what were some memorable aspects of that?

Well, that was the height of the dot com boom bubble bust. I mean, that was the, that was the frenzy. If you watch a succession, that was the, that was like what they did in the fourth season with the, with Gojo and the merger with the, with their, so it was like the height of like one of those kind of dynamics. And so

Would you recommend succession, by the way? I'm more of a Yellowstone guy.

Yellowstone's very American. I'm, I'm very proud of you. That's, that is I just talked to Matthew McConaughey and I'm full on Texan at this point.

Good. I hurtily approve. And he will be doing the sequel to Yellowstone.

Yeah. Very, very exciting. Anyway,

so that's a rude interruption by me by way of succession.

So that was at the height of the deal making and money and just the fur flying and like craziness.

And so, yeah, it was just one of those. It was just like, I mean, as the entire Netscape thing from start to finish was four years, which was like for, for one of these companies, it's just like incredibly fast. You know, we went public 18 months after we got, after we were founded, which virtually never happens. So it was just this incredibly fast kind of meteor streaking across the sky. And then of course it was this, and then there was just this explosion, right? That happened because then it was almost immediately followed by the dot com crash. It was then followed by Haywell buying Time Warner, which again is the succession guys going to play with that, which turned out to be a disastrous deal. You know, one of the famous, you know, kind of disasters in business history. And then, and then, you know, what became an internet depression on the other side of that. But then in that depression in the 2000s was the beginning of broadband and smartphones and web 2.0, right? And then social media and search and every SaaS and everything that came out of that. So what did you learn from just the acquisition? I mean, this is so much money. What's interesting, because I must have been very new to you, that the software stuff, you can make so much money. There's so much money swimming around. I mean, I'm sure the ideas of investment was starting to get born there.

Yes. Let me get, so let me lay out. So here's, here's the thing I don't know if I figured out them, but figured out later, which is software is a technology that it's like, you know, the concept of the philosopher stone, the philosopher stone in Alchemy transmutes lead into gold and Newton spent 20 years trying to find the philosopher stone, never got there. Nobody's ever figured it out. Software is our modern philosopher stone. And in economic terms, it transmutes labor into capital, which is like a super interesting thing. And by the way, like Karl Marx is rolling over in his grave right now, because of course, that's a complete refutation of his entire theory. Transmute labor into capital, which is, which is as follows is somebody sits down at a keyboard and types a bunch of stuff in and a capital asset comes out the other side. And then somebody buys that capital asset for a billion dollars. Like, that's amazing. Right. It's literally creating value right out of thin air, right? Out of purely human thought, right? And so that's, that's, there are many things that make software magical and special, but that's the economics. I wonder what Marx would have thought about that. Oh, he would have completely broke his brain because of course, the whole, the whole thing was, he was, you could, he, you know, that kind of technology is inconceivable when he was alive. It was all, it's all industrial era stuff. And so the, any kind of machinery necessarily involved huge amounts of capital and then labor was on the, on the, on the receiving end of the abuse, right? But like software, software, a software engineer is somebody who basically transmits his own labor into action, an actual capital asset creates permanent value. Well, in fact, it's actually very inspiring. That's actually more true today than before. So when, when I was doing software, the assumption was all new software basically has a sort of a parabolic sort of lifecycle, right? So you, you ship the thing, people buy it. At some point, everybody who wants it has bought it and then it becomes obsolete and it's like bananas. Nobody, nobody buys old software. These days, Minecraft, Mathematica, you know, Facebook, Google, you have the software assets that are, you know, have been around for 30 years that are gaining in value every year, right? And they're just, they're being World of Warcraft, right? Salesforce.com, like they're being, every single year they're being polished and polished and polished and polished,

they're getting better and better, more powerful, more powerful, more valuable, more valuable. So we've entered this era where you can actually have these things that actually build out over decades, which by the way, is what's happening right now with GPT. And so now, and this is why, you know, there, there, there is always, you know, sort of a constant investment frenzy around software is because, you know, look, when you start one of these things, it doesn't always succeed. But when it does, now you might be building an asset that builds value for, you know, four or five, six decades to come. You know, if you have a team of people who have the level of devotion required to keep making it better. And then the fact that, of course, everybody's online, you know, there's five billion people that are a click away from any new pieces software. So the potential market size for any of these things is, you know, nearly infinite. There must have been surreal back then, though. Yeah, yeah. This was all brand new, right? Yeah. Back then, this was all brand new. These were all, you know, brand new. Had you rolled out that theory in even 1999, people would have thought you were spucking crack. So that's, that's emerged over time. Well, let's now turn back into the future. You wrote the essay, why AI will save the world. Let's start at a very high level. What's the main thesis of the essay? Yeah. So the main thesis on the essay is that what we're dealing with here is intelligence. And it's really important to kind of talk about the sort of very nature of what intelligence is. And fortunately, we have a, we have a predecessor to machine intelligence, which is human intelligence. And we've got, you know, observations and theories over thousands of years for what intelligence is in the hands of humans and what intelligence is, right? I mean, what it literally is, is the way to, you know, capture, process, analyze, synthesize, information, solve problems. But the observation of intelligence in human hands is that intelligence quite literally makes everything better. And what I mean by that is every kind of outcome of like human quality of life, whether it's education outcomes or success of your children or career success or health or lifetime satisfaction. By the way, propensity to peacefulness as opposed to violence, propensity for open-mindedness versus bigotry, those are all associated with higher levels of intelligence. Smarter people have better outcomes than almost as you write in almost every domain of activity, academic achievement, job performance, occupational status, income, creativity, physical health, longevity, learning new skills, managing complex tasks, leadership, entrepreneurial success, conflict resolution, reading comprehension, financial decision making, understanding other perspectives, creative arts, parenting outcomes and life satisfaction. One of the more depressing conversations I've had, and I don't know why it's depressing. I have to really think through why it's depressing. But on IQ and the G factor, and that that's something in large part is genetic and it correlates so much with all of these things and success in life. It's like all the inspirational stuff we read about, like if you work hard and so on, damn, it sucks that you're born with the hand that you can't change. But what if you could? You're saying basically, a really important point, and I think it's in your articles, it really helped me. It's a nice added perspective to think about, listen, human intelligence, the science of intelligence has shown scientifically that it just makes life easier and better, the smarter you are. And now let's look at artificial intelligence. And if that's a way to increase the some human intelligence, then it's only going to make a better life. That's the argument. And certainly at the collective level, we could talk about the collective effect of just having

more intelligence in the world, which will have very big payoff. But there's also just at the individual level, like what if every person has a machine, you know, and it's a concept of argument, Doug Engelbar's concept of augmentation, you know, what if everybody has an assistant, and the assistant is, you know, 140 IQ, and you happen to be 110 IQ, and you've got, you know, something that basically is infinitely patient and knows everything about you and is pulling for you in every possible way, wants you to be successful. And anytime you find anything confusing, or want to learn anything, or have trouble understanding something, or want to figure out what to do in a situation, right, want to figure out how to prepare for a job interview, like any of these things, like it will help you do it. And it will therefore, the combination will effectively be, you know, effectively raise your raise, because it will effectively raise your IQ will therefore raise the odds of successful life outcomes in all these areas.

So people below the this hypothetical 140 IQ, you'll pull them off towards 140 IQ.

Yeah, yeah. And then of course, you know, people at people at 140 IQ will be able to have a peer, right, to be able to communicate, which is great. And then people above 140 IQ will have an assistant

that they can farm things out to. And then look, got willing, you know, at some point, these things go from future versions go from 140 IQ equivalent to 150 to 160 to 180, right, like Einstein was estimated to be on the order of 160. You know, so when we get, you know, 160 AI, like will be, you know, one assumes creating Einstein level breakthroughs and physics. And and then at 180 will be, you know, carrying cancer and developing warp drive and doing all kinds of stuff. And so it is quite possibly the case, this is the most important thing that's ever happened, the best thing that's ever happened, because precisely because it's a lever on this single fundamental factor of intelligence, which is the thing that drives so much of everything else. Can you still man the case that human plus AI is not always better than human for the individual?

You may have noticed that there's a lot of smart assholes running around. Sure. Yes. Right. And so like it's smart. There are certain people where they get smarter, you know, they get to be more arrogant, right? So, you know, there's one huge flaw. Although to push back on that, it might be interesting because when the intelligence is not all coming from you, but from a system, from another system that might actually increase the amount of humility, even in the assholes. One would hope. Yeah. Or it could make assholes more asshole. You know, that's, I mean, that's, that's for psychology to study. Yeah, exactly. Another one is smart people are very convinced that they, you know, have a more rational view of the world and that they have a easier time seeing through conspiracy theories and hoaxes and right, you know, sort of crazy beliefs and all that. There's a theory in psychology, which is actually smart people. So for sure, people who aren't as smart are very susceptible to hoaxes and conspiracy theories. But it may also be the case that the smarter you get, you become susceptible in a different way, which is you become very good at marshaling facts to fit preconceptions, right? You become very, very good at assembling whatever theories and frameworks and pieces of data and graphs and charts you need to validate whatever crazy ideas got in your head. And so you're susceptible in a different way, right? We're all sheep, but different colored sheep. Some sheep are better at justifying it, right? And those are, you know, those are the smart sheep, right? So yeah, look, like, I would say this, look, like there are no panacea. I'm not, I'm not a utopian. There are no

panaceas in life. There are no, like, you know, I don't believe they're like pure positives. I'm not a transcendental kind of person like that. But, you know, so yeah, there are going to be issues. And, and, you know, look smart people, maybe you could say about smart people as they are more likely to get themselves in situations that are, you know, beyond their grasp, you know, because they're just more confident in their ability to deal with complexity and their, their eyes become bigger, their cognitive eyes become bigger than their stomach. You know, so yeah, you could argue those eight different ways. Nevertheless, on that, right, clearly overwhelmingly, again, if you just extrapolate from what we know about human intelligence, you're improving so many aspects of life if you're upgrading intelligence. So there'll be assistance at all stages of life. So when you're younger, there's for education, all that kind of stuff, for mentorship, all of this. And later on, as you're doing work and you've developed a skill and you're having a profession, you'll have an assistant that helps you excel at that profession. So at all stages of life. Yeah, I mean, look, the theory is augmentations. This is the DeGengelbert's term for it. DeGengelbert made this observation many, many decades ago that, you know, basically, it's like you can have this oppositional frame of technology where it's like us versus the machines. But what you really do is you use technology to augment human capabilities. And by the way, that's how actually the economy develops. That's the economic side of this, but that that's actually how the economy grows is through technology augmenting human human potential. And so yeah, and then you basically have a proxy or a, you know, or a, you know, a sort of prosthetic, you know, so like you've got glasses, you've got a wristwatch, you know, you've got shoes, you know, you've got these things, you've got a personal computer, you've got a word processor, you've got Mathematica, you've got Google. This is the latest viewed through that lens. AI is the latest in a long series of basically augmentation methods to be able to raise human capabilities. It's just this one is the most powerful one of all, because this is the one that goes directly to what they call fluid intelligence, which is IQ. Well, there's two categories of folks that you outline that they worry about or highlight the risks of AI. And you highlight a bunch of different risks. I'd love to go through those risks and just discuss them brainstorm, which ones are serious and which ones are less serious. But first, the Baptist and the bootleggers, what are these two interesting groups of folks who are, who, who worry about the effect of AI on human civilization? Or say they do.

The Baptist worry the bootleggers say they do. So the Baptist and the bootleggers is a metaphor from economics from what's called development economics. And it's this observation that when you get social reform movements in a society, you tend to get two sets of people showing up arguing for the social reform. And the term Baptist and bootleggers comes from the American experience with alcohol prohibition. And so in the 1900s, 1910s, there was this movement that was very passionate at the time, which basically said alcohol is evil, and it's destroying society. By the way, there was a lot of evidence to support this. There were very high rates of, very high correlations than, by the way, and now between rates of physical violence and alcohol use, almost all violent crimes have either the perpetrator or the victim are both drunk. You see this actually in the work, almost all sexual harassment cases in the workplace. It's like at a company party and somebody's drunk. It's amazing how often alcohol actually correlates to actually dysfunction, at least to domestic abuse and so forth, child abuse. And so you had this group of people who were like, okay, this is bad stuff and we shall outlaw it. And those were

quite literally the Baptists. Those were super committed, hardcore Christian activists in a lot of cases. There was this woman whose name was Carrie Nation, who was this older woman who had

been in this, I don't know, disastrous marriage or something, and her husband had been abusive and drunk all the time. And she became the icon of the Baptist prohibitionists. And she was legendary in that era for carrying an axe and doing, completely on her own, doing raids of saloons and taking her axe to all the bottles and tigs in the back. So a true believer.

An absolute true believer with absolutely the purest of intentions. And again, there's a very important thing here, which is you could look at this cynically and you could say the Baptists are like delusional, you know, extremists, but you can also say, look, they're right. Like she was, you know, she had a point. Like she wasn't wrong about a lot of what she said. But it turns out the way the story goes is it turns out that there were another set of people who very badly wanted to outlaw alcohol in those days. And those were the bootleggers, which was organized crime that stood to make a huge amount of money if legal alcohol sales were banned. And this was in fact, the way the history goes is this was actually the beginning of organized crime in the US. This was the big economic opportunity that opened that up. And so they went in together. And they didn't go in together. Like the Baptists did not even necessarily know about the bootleggers because they were on the moral crusade. The bootleggers certainly knew about the Baptists. And they were like, wow, this is these people are like the great front people for like, you know, it's good PR shenanigans in the background. And they got the full state act passed. Right. And they did in fact ban alcohol in the US. And you'll notice what happened, which is people kept drinking. It didn't work. People kept drinking that bootleggers made a tremendous amount of money. And then over time, it became clear that it made no sense to make it illegal and it was causing more problems. And so then it was revoked. And here we sit with legal alcohol 100 years later with all the same problems. And, you know, the whole thing was this like giant misadventure. The Baptist got taken advantage of by the bootleggers and the bootleggers got what they wanted. And that was that the same two categories of folks are now sort of suggesting that the development of artificial intelligence should be regulated.

100%. Yeah, it's the same pattern. And the economists will tell you it's the same pattern every time. Like this is what happened with nuclear power. This is what happened, which is another interesting one. But like, yeah, this is this happens dozens and dozens of times throughout the last 100 years. And this is what's happening now.

And you write that it isn't sufficient to simply identify the actors and impugn their motors. We should consider the arguments of both the Baptists and the bootleggers on their merits. So let's do just that. Risk number one.

Will AI kill us all? Yes.

So what do you think about this one? What do you think is the core argument here that the development of AGI, perhaps better said, will destroy human civilization?

Well, first of all, you just did a slight of hand because we went from talking about AI to AGI. Is there a fundamental difference there? I don't know. What's AGI?

What's AI? What's intelligence? Well, I know what AI is. AI is machine learning.

What's what's AGI? I think we don't know what the bottom of the well of machine learning is or what the ceiling is. Because just to call something machine learning or

just to call something statistics or just to call it math or computation doesn't mean, you know, nuclear weapons are just physics. So it's, to me, it's very interesting and surprising how far machine learning has taken. No, but we knew that nuclear physics would lead to weapons. That's why the scientists of that era were always in some of this huge dispute about building the weapons. This is different. AGI is different. Where does machine learning lead? Do we know? We don't know, but this, my point is different. We actually don't know. And this is where you, the slight of hand kicks in, right? This is where it goes from being a scientific topic to being a religious topic. And that's why I specifically called out the, because that's what happens. They do the vocabulary shift. And all of a sudden, you're talking about something totally that's not actually real.

Well, then maybe you can also, as part of that, define the Western tradition of millennialism.

Yes. Into the world. Apocalypse. Apocalypse cults.

Apocalypse cults.

Well, so we live in, we, of course, live in a Judeo-Christian, but primarily Christian, kind of saturated, you know, kind of Christian, post-Christian, secularized Christian, you know, kind of world in the West. And of course, court of Christianity is the idea of the second coming and, you know, the revelations and, you know, Jesus returning in the thousand year, you know, utopia on earth and then the, you know, the rapture and like all that stuff.

You know, we don't, we, you know, we collectively, you know, as a society, we don't necessarily take all that fully seriously now. So what we do is we create our secularized versions of that. We keep, we keep looking for utopia. We keep looking for, you know, basically the end of the world. And so what you see over, over decades is that basically a pattern of these sort of, of these, of these, of this, this is what cults are. This is how cults form as they form around some theory of the end of the world. And so the people's temple cult, the Manson cult, the heaven's gate cult, the David Koresh cult, you know, what they're all organized around is like, there's going to be this thing that's going to happen that's going to basically bring civilization crashing down. And then we have this special elite group of people who are going to see it coming and prepare for it. And then they're the people who are either going to stop it or a failing stopping it. They're going to be the people who survive to the other side and ultimately get credit for having been right. Why is that so compelling? Do you think? Because it satisfies this very deep need we have for transcendence and meaning that got stripped away when we became secular. Yeah, but why is the transcendence involved, the destruction of human civilization? Because like, how plausible, it's like a very deep psychological thing because it's like, how plausible, how plausible is it that we live in a world where everything's just kind of all right? Right? How exciting is that?

We want more than that. But that's the deep question I'm asking.

Why is it not exciting to live in a world where everything's just all right? Because I think most of the animal kingdom would be so happy with just all right. Because that means survival. Maybe that's what it is. Why are we conjuring up things to worry about?

So CS Lewis called it the God-shaped hole. So there's a God-shaped hole in the human experience, consciousness, soul, whatever you want to call it, where there's got to be something that's bigger than all this. There's got to be something transcendent. There's got to be something that is bigger, right? Bigger, bigger purpose of bigger meaning. And so we have run the experiment

of, you know, we're just going to use science and rationality and kind of, you know, everything's just going to kind of be as it appears. And a large number of people have found that very deeply wanting and have constructed narratives. And by the way, this is the story of the 20th century, right? Communism was one of those. Communism was a form of this. Nazism was a form of this. You know, some people, you know, you can see movements like this playing out all over the world right now. So you construct a kind of devil, a kind of source of evil, and we're going to transcend beyond it. Yeah. And the millenarian, the millenarian is kind of, when you see a millenarian cult, they put a really specific point on it, which is end of the world, right? There is some change coming. And that change that's coming is so profound and so important that it's either going to lead to utopia or hell on earth, right? And it is going to, and then, you know, it's like, what if you actually knew that that was going to happen, right? What would you, what would you do, right? How would you prepare yourself for it? How would you come together

with a group of like-minded people, right? How would you, what would you do? Would you plan like ashes of weapons in the woods? Would you like, you know, I don't know, create underground buckers? Would you, you know, spend your life trying to figure out a way to avoid having it happen? Yeah, that's a really compelling, exciting idea to, to have a club over, to have a, to have a, to have a little bit of trouble. Like you get together on a Saturday night and drink some beers and talk about the, the end of the world and how you, you are the only ones who have figured it out. Yeah. And then, and then once you lock in on that, like, how can you do anything else with your life? Like this is obviously the thing that you have to do. And then, and then there's a psychological effect you alluded to. There's a psychological effect. If you take a set of true believers and you leave them to themselves, they get more radical, right? Because they, they self radicalize each other. That said, it doesn't mean they're not sometimes right. Yeah, the end of the world might be, yes, correct. Like, they might be right. Yeah.

But like, we have some pamphlets for you. I mean, there's, I mean, we'll talk about nuclear weapons because you have a really interesting little moment that I learned about in your essay. But, you know, sometimes it could be right. Yeah. Because we're still, you were developing more and more powerful technologies in this case. And we don't know what the impact they will have on human civilization. Well, we can highlight all the different predictions about how it will be positive. But the risks are there. And you discuss some of them. Well, the steel man, the steel man is the steel man, actually, the steel man and his refutation are the same, which is, well, you can't predict what's going to happen, right? You, right? You can't rule out that this will not end everything, right? But the response to that is you have just made a completely non-scientific claim. You've made a religious claim, not a scientific claim. How does it get disproven?

There is, and there's no, by definition, with these kinds of claims, there's no way to disprove them, right? And so there's no, you just go right on the list. There's no hypothesis. There's no testability of the hypothesis. There's no way to falsify the hypothesis. There's no way to measure progress along the arc. Like, it's just all completely missing. And so it's not scientific.

Well, I don't think it's completely missing. It's somewhat missing. So, for example, the people that say, yeah, it's going to kill all of us. I mean, they usually have ideas about how to do that, whether it's the paperclip maximizer or, you know, it escapes. There's mechanism by which you can imagine it killing all humans. And you can disprove it by saying there is a limit to

the speed at which intelligence increases. Maybe show that, like, they sort of rigorously really describe model, like how it could happen and say, no, here's a physics limitation. There's a physical limitation to how these systems would actually do damage to human civilization. And it is possible they will kill 10 to 20% of the population, but it seems impossible for them to kill 99%. It was practical counterarguments, right? So you mentioned basically what I described as the thermodynamic counterargument, which is sitting here today. It's like we're with the evil AGI get the GPUs because they don't exist. So you're going to have a very frustrated baby evil AGI who's going to be trying to buy NVIDIA stock or something to get them to finally make some chips. So the serious form of that is the thermodynamic argument, which is like, okay, where's the energy going to come from? Where's the processor going to be running? Where's the data center going to be happening? How is this going to be happening in secret such that you know it's not, you know? So that's a practical counterargument to the runaway AGI thing. But we can argue that and discuss that. I have a deeper objection to it, which is this is all forecasting. It's all modeling. It's all future prediction. It's all future hypothesizing. It's not science. It is the opposite of science. So they'll pull up Carl Sagan, extraordinary claims require extraordinary proof, right? These are extraordinary claims. The policies that are being called for to prevent this are of extraordinary magnitude. And I think we're going to cause extraordinary damage. And this is all being done on the basis of something that is literally not scientific. It's not a testable hypothesis. So the moment you say AI is going to kill all of us, therefore, we should ban it or that we should regulate all that kind of stuff, that's when it starts getting serious. Or start, you know, military airstrikes and data centers. Oh boy. Right? And like, yeah, this one gets starts. Well, so it starts getting real. So here's the problem of millinery and cults. They have a hard time staying away from violence. But violence is so fun. If you're on the right end of it, they have a hard time avoiding violence. The reason they have a hard time avoiding violence is if you actually believe the claim, right, then what would you do to stop the end of the world? Well, you would do anything, right? And so and this is where you get and again, if you just look at the history of millinery and cults, this is where you get the people's temple and everybody killing themselves in the jungle. And this is where you get Charles Manson and you know, sending in me to kill kill the pigs. Like this is the problem with these. They have a very hard time drawing the line at actual violence. And I think I think in this case, there's, I mean, they're already calling for it like today. And you know, where this goes from here as they get more worked up, like I think it's like really concerning. Okay. But that's kind of the extremes, you know, the extremes of anything I was concerning. It's also possible to kind of believe that AI has a very high likelihood of killing all of us. But there's and therefore we should maybe consider slowing development or regulating. So not violence or any of these kinds of things, but saying like, all right, let's, let's take a pause here. You know, biological weapons, nuclear weapons, like whoa, this is like serious stuff. We should be careful. So it is possible to kind of have a more rational response, right? If you believe this risk is real. Believe. Yes. So is it possible to be have a scientific approach to the prediction of the future? I mean, we just went through this with COVID. What do we know about modeling? Well, I mean, what do we learn about modeling with COVID? There's a lot of lessons. They didn't work at all.

They worked poorly. The models were terrible. The models were useless. I don't know if the models were useless or the people interpreting the models and then the centralized institutions that were creating policy rapidly based on the models and leveraging the models in order to support their narratives versus actually interpreting the airbars and the models and all that kind of stuff. What you had with COVID, in my view, you had with COVID is you had these experts showing up and they claimed to be scientists and they had no testable hypotheses whatsoever. They had a bunch of models. They had a bunch of forecasts and they had a bunch of theories and they laid these out in front of policymakers and policymakers freaked out and panicked, right? And implemented a whole bunch of like really like terrible decisions that we're still living with the consequences of. And there was never any empirical foundation to any of the models. None of them ever came true. Yeah, to push back, there were certainly Baptist and bootleggers in the context of this pandemic, but there's still a usefulness to models, no? So not if they're reliably wrong, right? Then they're actually like anti-useful, right? They're actually damaging. But what do you do with the pandemic? What do you do with any kind of threat? Don't you want to kind of have several models to play with as part of the discussion of like, what the hell do we do here? I mean, do they work? Because they're an expectation that they actually like work, that they have actual predictive value? I mean, as far as I can tell with COVID, the policymakers just sigh out themselves into believing that there was some, I mean, look, the scientists, the scientists were at fault. The quote on quote, scientists showed up. So I had to submit that into this. So there was a, remember the Imperial College models out of London were the ones that were like, these are the gold standard models. So a friend of mine runs a big software company and he was like, wow, this is like COVID's really scary. And he's like, you know, he contacted this research and he's like, you know, do you need some help? You've been just building this model on your own for 20 years. Do you need some food? Do you like our coders to basically restructure it so it can be fully adapted for COVID? And the guy said yes and sent over the code. And my friend said it was like the worst spaghetti code he's ever seen. That doesn't mean it's not possible to construct a good model of pandemic with the correct error bars with a high number of parameters that are continuously many times a day updated as we get more data about a pandemic. I would like to believe when a pandemic hits the world, the best computer scientists in the world, the best software engineers respond aggressively. And as input take the data that we know about the virus and it's an output, say here's, here's what's happening in terms of how quickly it's spreading, what that lead in terms of hospitalization and death and all that kind of stuff. Here's how likely how contagious it likely is. Here's how deadly it likely is based on different conditions, based on different ages and demographics and all that kind of stuff. So here's the best kinds of policy. It feels like you could have models, machine learning, that kind of, they don't perfectly predict the future, but they help you do something because there's pandemics that are like meh, they don't really do much harm. And there's pandemics, you can imagine them, they could do a huge amount of harm, like they can kill a lot of people. So you should probably have some kind of data driven models that keep updating, that allow you to make decisions that are basically like, where, how bad is this thing?

Now you can criticize how horrible all that went with the response to this pandemic, but I just feel like there might be some value to models. So to be useful at some point, it has to be predictive, right? So the easy thing for me to do is to say, obviously, you're right. Obviously, I want to see that just as much as you do, because anything that makes it easier to navigate through society through a wrenching risk like that, that sounds great. The harder objection to it is just simply you are trying to model a complex dynamic system with 8 billion moving parts, like not possible.

It's very tough.

Can't be done, complex systems can't be done.

Machine learning says hold my beer, but is it possible? No?

I don't know. I would like to believe that it is. I'll put it this way. I think where you and I would agree is I think we would like that to be the case. We are strongly in favor of it. I think we would also agree that respect to COVID or pandemics, no such thing, at least neither you nor I think are aware. I'm not aware of anything like that today.

My main worry with the response to the pandemic is that, same as with aliens, is that even if such a thing existed, and as possible it existed, the policymakers were not paying attention.

There was no mechanism that allowed those kinds of models to percolate out.

I think we had the opposite problem during COVID. I think the policymakers,

I think these people with basically fake science had too much access to the policymakers.

Right, but the policymakers also wanted, they had a narrative in mind and they also wanted to use whatever model that fit that narrative to help them out. It felt like there was a lot of politics and not enough science. Although a big part of what was happening, a big reason we got lockdowns for as long as we did was because these scientists came in with these doomsday scenarios that were just completely off the hook. Scientists and quotes, that's not quote unquote scientists. That's give love. So here's science. That is the way out.

Science is a process of testing hypotheses. Modeling does not involve testable hypotheses, right? I don't even know that modeling actually qualifies to science.

Maybe that's a side conversation we could have some time over a beer.

That's really interesting, but what do we do about the future? I mean, what?

So number one is when we start with number one, humility,

goes back to this thing of how do we determine the truth. Number two is we don't believe, I've got a hammer, everything looks like a nail, right? This is one of the reasons I gave you, I gave Lex a book, which the topic of the book is what happens when scientists basically stray off the path of technical knowledge and start to weigh in on politics and societal issues.

In this case, philosophers.

Well, in this case, philosophers. But he actually talks in this book about Einstein.

He talks about the nuclear age in Einstein. He talks about the physicists actually doing very similar things at the time. The book is One Reason Goes on Holiday Philosophers and Politics by Nevin. And it's just a story. It's a story. There are other books on this topic, but this is a new one that's really good. It's just a story of what happens when experts in a certain domain decide to weigh in and become basically social engineers and basically political advisers. And it's just a story of just unending catastrophe, right? And I think that's what happened with COVID

again. Yeah, I found this book a highly entertaining and eye-opening read filled with amazing anecdotes of irrationality and craziness by famous recent philosophers.

After you read this book, you will not look at Einstein the same.

Oh boy. Yeah.

Don't destroy my heroes.

You will not be a hero of yours anymore.

I'm sorry. You probably shouldn't read the book.

All right.

But here's the thing. The AI risk people, they don't even have the COVID model.

At least not that I'm aware of.

No.

Like there's not even the equivalent of the COVID model. They don't even have the spaghetti code.

They've got a theory and a warning and a this and a that. And like if you ask like, okay,

well, here's the ultimate example is, okay, how do we know, right? How do we know that an AI is running away? Like how do we know that the FOOM takeoff thing is actually happening?

And the only answer that any of these guys have given that I've ever seen is, oh,

it's when the loss rate, the loss function in the training drops, right?

That's when you need to like shut down the data center, right?

And it's like, well, that's also what happens when you're successfully training a model.

Like, what even is, this is not science. This is not, it's not anything. It's not a model.

It's not anything. There's nothing to arguing with it. It's like, you know, punching Jello,

like there's, what do you even respond to?

So just put pushback on that. I don't think they have good metrics of, yeah, when the FOOM is happening, but I think it's possible to have that. Like I just, just as you speak now, I mean, it's possible to imagine there could be measures.

It's been 20 years.

No, for sure. But it's been only weeks since we had a big enough breakthrough in language models. We can start to actually have, this, the thing is, the AI do more stuff

didn't have any actual systems to really work with it. Now there's real systems,

you can start to analyze like, how does this stuff go wrong? And I think you kind of

agree that there is a lot of risks that we can analyze. The benefits outweigh the risks in many

cases. Well, the risks are not existential. Yes. Well, not in the, not in the FOOM, not in the FOOM

paperclip, not this. Let me, okay, there's another slide of hand that you just alluded to. There's

another slide of hand that happens, which is very... I think I'm very good at the slide of hand thing.

Not scientific. So the book Superintelligence, right, which is like the Nick Bostrom's book,

which is like the origin of a lot of this stuff, which was written, you know, whatever,

10 years ago or something. So he does this really fascinating thing in the book, which is he basically

says, there are many possible routes to machine intelligence, to artificial intelligence, and

he describes all the different routes to artificial intelligence, all the different

possible, everything from biological augmentation through to, you know, that all these different

things. One of the ones that he does not describe is large language models, because of course,

the book was written before they were invented, and so they didn't exist. In the book, he describes

them all, and then he proceeds to treat them all as if they're exactly the same thing. He presents

them all as sort of an equivalent risk to be dealt with in an equivalent way to be thought about the same way. And then the risk, the quote unquote risk that's actually emerged is actually a completely different technology than he was even imagining. And yet all of his theories and beliefs are being transplanted by this movement, like straight onto this new technology. And so again, like, there's no other area of science or technology where you do that. Like when you're dealing with like organic chemistry versus inorganic chemistry, you don't just like say, oh, with respect

to like either one, basically, maybe, you know, growing up and eating the world or something, like they're just going to operate the same way, like you don't. But you can start talking about, like as we get more and more actual systems that start to get more and more intelligent, you can start to actually have more scientific arguments here. Like, you know, high level, you can talk about the threat of autonomous weapons systems back before we had any automation in the military. And that would be like very fuzzy kind of logic. But the more and more you have drones, they're becoming more and more autonomous, you can start imagining, okay, what does that actually look like? And what's the actual threat of autonomous weapons systems? How does it go wrong? And still it's, it's, it's very vague, but you start to get a sense of like, all right, it should probably be illegal or wrong or not allowed to do like mass deployment of fully autonomous drones that are doing aerial strikes on large areas. I think it should be required. Right. So that's, no, no, no, no, I think it should be required that only aerial vehicles are automated. Okay, so you want to go the other way? I want to go the other way. So that, okay, I think it's obvious that the machine is going to make a better decision than the human pilot. I think it's obvious that it's in the best interest of both the attacker and the defender and humanity at large, if machines are making more decisions than not people. I think people make terrible decisions in times of war. But like there's a, there's ways this can go wrong too, right? Well, it's worse go terribly wrong now. This goes back to the, this is that whole thing about like the self-driving, does the self-driving car need to be perfect versus does it need to be better than the human driver? Yeah. Does the automated drone need to be perfect or does it need to be better than a human pilot at making decisions under enormous amounts of stress and uncertainty? Yeah. Well, the, on average, the, the worry that AI folks have is the runaway. They're going to come alive, right? Then again, that's the sleight of hand, right? Or not, not come alive. Well, hold on a second. You become, you lose control as well. But then they're going to develop goals of their own. They're going to develop a mind of their own. They're going to develop their own, right? No, more, more like Chernobyl style meltdown, like just bugs in the code, accidentally, you know, force you the results in the bombing of like large civilian areas. Okay. To a degree that's not possible in the, in the current military strategies, patrol by humans. I don't know. Well, actually, we've been doing a lot of mass bombings of cities for a very long time. Yes. And a lot of civilians died. And a lot of civilians died. And if you watch the documentary, The Fog of War, McNamara, it spends a big part of it talking about the firebombing of the Japanese cities, burning them straight to the ground, right? The devastation in Japan, American military firebombing, the cities in Japan was a considerably bigger devastation than the use of nukes, right? So we've been doing that for a long time. We also did that to Germany, by the way, Germany did that to us, right? Like, that's an old tradition. The minute we got airplanes, we started doing indiscriminate bombing.

So one of the things that we're still doing it, the modern US military can do with technology, with automation, but technology more broadly is higher and higher precision strikes. Yeah. And so precision is obviously, and this is the, the JDAM, right? So there was this big advance, this is a big advance called the JDAM, which basically was strapping a GPS transceiver to a, to a, to an unguided bomb and turning it into a guided, guided bomb. And yeah, that's great. Like, look, that's been a big advance. But, and that's like a baby version of this question, which is, okay, do you want like the human pilot, like guessing where the bomb's going to land, or do you want like the machine, like guiding the bomb to its destination? That's a baby version of the question. The next version of the question is, do you want the human or the machine deciding whether to drop the bomb? Everybody just assumes the human's going to do a better job for what I think are fundamentally suspicious reasons. Emotional psychological reasons. I think it's very clear that the machine's going to do a better job making that decision, because the human's making that, making that decision or God awful, just terrible. Yeah. Right. And so, so yeah, so this is the, this is the thing. And then let's get to the, there was, can I, one more slide of hand? Yes, sure. Okay. I'm a magician, you could say. One more slide of hand. These things are going to be so smart, right? That they're going to be able to destroy the world and wreak havoc and like do all this stuff and plan and do all this stuff and evade us and have all their secret things and their secret factories and all this stuff. But they're so stupid that they're going to get like tangled up in their code. And that's the thing. They're not going to come alive, but there's going to be some bug that's going to cause them to like turn us all in a paper. Like that they're not going to, that they're going to be genius in every way other than the actual bad goal. And it's just like, and that's just like a like ridiculous like discrepancy. And, and, and, and you can prove this today. You can actually address this today for the first time with LMS, which is you can actually ask LMS to resolve moral dilemmas. Yeah. So you can create the scenario, you know, dot, dot, dot, this, that, this, that, this, that. What would you as the AI do in the circumstance? And they don't just say destroy all humans, destroy all humans. They will give you actually very nuanced moral, practical, trade off oriented answers. And so we actually already have the kind of AI that can actually like think this through and can actually like, you know, reason about goals. Well, the hope is that AGI or like a very super intelligent systems have some of the nuance that LMS have. And the intuition is they most likely will because even these LMS have the nuance. LMS are really, this is actually worth, worth spending a moment on LMS are really interesting to have moral conversations with. And that, I didn't expect I'd be having a moral conversation with a machine in my lifetime. And let's remember, we're not really having a conversation with a machine where we're having a conversation with the entirety of the collective intelligence of the human species. Exactly. Yeah, correct. But it's possible to imagine autonomous weapons systems that are not using LMS. But if they're smart enough to be scary, where are they not smart enough to be wise? Like that's the part where it's like, I don't know how you get the one without the other. Is it possible to be super intelligent without being super wise? Well, you're, again, you're back to, I mean, then you're back to a classic autistic computer, right? Like you're back to just like a blind rule follower. I've got this like core is the paperclip thing, I've got this core rule, and I'm just going to follow it to the end of the earth. And it's like, well, but everything

you're going to be doing to execute that rule is going to be super genius level that humans aren't going to be able to counter. It's just, it's a mismatch in the definition of what the system is capable of. Unlikely, but not impossible, I think. But again, here you get to like, okay, like, no, I'm not saying when it's unlikely, but not impossible. If it's unlikely, that means the fear should be correctly calibrated. Extraordinary claims require extraordinary proof.

Well, okay. So one interesting sort of tangent, I would love to take on this because you mentioned this in the essay about nuclear, which was also, I mean, you don't shy away from a little bit of a spicy take. So Robert Oppenheimer famously said, now I am become death the destroyer of worlds as he witnessed the first detonation of a nuclear weapon on July 16, 1945. And you write an interesting historical perspective, quote, recall that John von Neumann responded to Robert Oppenheimer's famous hand wringing about the role of creating nuclear weapons, which, you note, helped

end World War Two and prevent World War Three with some people confess guilt to claim credit for the sin. And you also mentioned that Truman was harsher after meeting Oppenheimer. He said that don't let that crybaby in here again. Real quote, by the way, from Dean Acheson.

Boy. Because Oppenheimer didn't just say the famous line. He then spent years going around basically moaning and going on TV and going into the White House and basically like just like doing this hair shirt thing, this sort of self-critical like, oh my God, I can't believe how awful I am.

So he's widely considered perhaps because of the hand wringing as the father of the atomic bomb. This is Von Neumann's criticism of him as he tried to have his cake and eat it too, like he wanted to. And so, Von Neumann, of course, is a very different kind of personality and he's just like, you know, this is like an incredibly useful thing. I'm glad we did it.

Yeah. Well, Von Neumann is as widely credited as being one of the smartest humans of the 20th century. Certain people, everybody says like, this is the smartest person I've ever met when they've met him. Anyway, that doesn't mean smart, doesn't mean wise. So, I would love to sort of, can you make the case both for and against the critique of Oppenheimer here? Because we're talking about nuclear weapons. Boy, do they seem dangerous.

Well, so the critique goes deeper and I left this out. Here's the real substance I left it out because I didn't want to dwell on nukes in my AI paper. But here's the deeper thing that happened. And I'm really curious, this movie coming out this summer, I'm really curious to see how far he pushes this because this is the real drama in the story, which is it wasn't just a question of our nukes good or bad, it was a question of should Russia also have them. And what actually happened

was Russia got the bomb, America invented the bomb, Russia got the bomb, they got the bomb through espionage, they got American and, you know, they got American scientists and foreign scientists working on the American project, some combination of the two basically gave the Russians the designs for the bomb. And that's how the Russians got the bomb. There's this dispute to this day of Oppenheimer's role in that. If you read all the histories, the kind of composite picture, and by the way, we now know a lot actually about Soviet espionage in that era, because there's been all this declassified material in the last 20 years that actually shows a lot of very interesting things. But if you're going to read all the histories, what you're going to get is Oppenheimer himself probably was not a, he probably did not hand over the nuclear secrets himself. However, he was close to many people who did, including family members. And there were other members

of the Manhattan Project who were Russian Soviet SS and did hand over the bomb. And so the view that Oppenheimer and people like him had that this thing is awful and terrible, and oh my God, and you know, all this stuff, you could argue fed into this ethos at the time that resulted in people thinking that the Baptist is thinking that the only principle thing to do is to give the Russians the bomb. And so the moral beliefs on this thing and the public discussion and the role that the inventors of this technology play, this is the point of this book, when they kind of take on this sort of public intellectual moral kind of thing, it can have real consequences, right? Because we live in a very different world today because Russia got the bomb than we would have lived in had they not gotten the bomb, right? The entire 20th century, second half of the 20th century would have played out very different had those people not given Russia the bomb. And so the stakes were very high then. The good news today is nobody sitting here today, I don't think worrying about like an analogous situation with respect to like, I'm not really worried that Sam Altman is going to decide to give the Chinese the design for AI, although he did just speak at a Chinese conference, which is interesting. But however, I don't think that's what's at play here. But what's at play here are all these other fundamental issues around what do we believe about this and then what laws and regulations and restrictions that we're going to put on it. And that's where I draw like a direct straight line. And anyway, and my reading of the history on nukes is like the people who were doing the full hair shirt public, this is awful, this is terrible, actually had like catastrophically bad results from taking those views. And that's what I'm worried is going to happen again. But is there a case to be made that you really need to wake the public up to the dangers of nuclear weapons when they were first dropped? Like really like educate them on like, this is extremely dangerous and destructive weapon. I think the education kind of happened quick and early. Like, how? It was pretty obvious. How? We dropped one bomb and destroyed an entire city.

Yeah, so 80,000 people dead. And look, but I don't like the reporting of that, you can report that in all kinds of ways. You can you can do all kinds of slants like war is horrible, war is terrible, you can do, you can make it seem like nuclear, the use of nuclear weapons is just a part of war and all that kind of stuff. Something about the reporting and the discussion of nuclear weapons resulted in us being terrified in awe of the power of nuclear weapons. And that potentially fed in a positive way towards the game theory of mutually assured destruction.

Well, so this gets to what actually happens. Some of us may be playing devil's advocate here. Yeah, sure, of course. Let's get to what actually happened and then kind of back into that. So what actually happened, I believe, and again, I think this is a reasonable reading of history is what actually happened was nukes then prevented World War Three. And they prevented World War Three

through the game theory of mutually assured destruction. Had nukes not existed, there would have been no reason why the Cold War did not go hot. And then the military planners at the time thought both on both sides thought that there was going to be World War Three on the plains of Europe and they thought there was going to be like 100 million people dead. It was the most obvious thing in the world to happen. And it's the dog that didn't bark. It may be the best single net thing that happened in the entire 20th century is that that didn't happen.

Yeah, actually, just on that point, you say a lot of really brilliant things. It hit me just as you were saying it. I don't know why it hit me for the first time, but we got two wars in a span of

like 20 years. Like we could have kept getting more and more World Wars and more and more ruthless.

It actually you could have had a US versus Russia war.

Yeah, you could have. By the way, there's another hypothetical scenario. The other hypothetical scenario is the Americans got the bomb, the Russians didn't, right? And then America is the big dog. And then maybe America would have had the capability to actually roll back the Iron Curtain. I don't know whether that would have happened, but like it's entirely possible, right? And the act of these people who had these moral positions about because they could forecast, they could model, they could forecast the future of how the technology would get used made a horrific mistake because they basically ensured that the Iron Curtain would continue for 50 years longer than it would have otherwise. And again, like these are counterfactuals. I don't know that that's what would have happened.

But like the decision to hand the bomb over was a big decision made by people who were very full of themselves. Yeah, but so me as an America, me as a person that loves America, I also wonder if US was the only ones with nuclear weapons. That was the argument for handing the bomb.

That was the guys who handed over the bomb. That was actually their moral argument.

I would probably not hand it over to, I would be careful about the regimes you hand it over to.

Maybe give it to like the British or something.

Like a democratically elected government. Well, there are people to this day who think that those Soviet spies did the right thing because they created a balance of terror as opposed to the US having just, and by the way, let me. Balance of terror. Let's tell the full version. Such a sexy ring to it. Okay, so the full version of the story is John von Neumann is a hero of both years in mind. The full version of the story is he advocated for a first strike. So when the US had the bomb and Russia did not, he advocated for, he said, we need to distract them right now.

Distract Russia. Yeah.

Yes.

Von Neumann.

Yes, because he said World War III is inevitable. He was very hardcore. His theory was World War III

is inevitable. We're definitely going to have a World War III. The only way to stop World War III is we have to take them out right now, and we have to take them out right now before they get the bomb because this is our last chance. Now, again, like.

Is this an example of philosophers and politics?

I don't know if that's in there or not, but this is in the standard bottom.

No, but it is.

Yeah, meaning is that.

Yeah, this is on the other side. So most of the case studies in books like this are the crazy people on the left. Von Neumann is a story arguably of the crazy people on the right.

Yes, stick to computing, John.

Well, this is the thing, and this is the general principle. It goes back to our core thing, which is like, I don't know whether any of these people should be making any of these calls because there's nothing in either Von Neumann's background or Oppenheimer's background or any of these people's background that qualifies them as moral authorities.

Yeah. Well, this actually brings up the point of NAI, who are the good people to reason about the morality, the ethics. Outside of these risks, outside of like the more complicated stuff that you agree on is this will go into the hands of bad guys and all the kinds of ways they'll do is interesting and dangerous, is dangerous in interesting and unpredictable ways, and who is the right person? Who are the right kinds of people to make decisions how to respond to it?

Are these the tech people?

So the history of these fields, this is what he talks about in the book, the history of these fields is that the competence and capability and intelligence and training and accomplishments of senior scientists and technologists working on a technology and then being able to then make moral judgments in the use of that technology, that track record is terrible. That track record is catastrophically bad. The people that develop that technology are usually not going to be the right people. So the claim is, of course, they're the knowledgeable ones, but the problem is they've spent their entire life in a lab. They're not theologians. So what you find when you read this, when you look at these histories, what you find is they generally are very thinly informed on history, on sociology, on theology, on morality, ethics. They tend to manufacture their own world views from scratch. They tend to be very thin. They're not remotely the arguments that you would be having if you got a group of highly qualified theologians or philosophers or...

Well, as the devil's advocate takes a sip of whiskey, say that I agree with that, but also it seems like the people who are doing the ethics departments and these tech companies go sometimes the other way. They're not nuanced on history or theology or this kind of stuff. It almost becomes a kind of outraged activism towards directions that don't seem to be grounded in history and humility and nuance. It's again drenched with arrogance. So I'm not sure which is worse.

Oh, no, they're both bad. Yeah, so definitely not them either.

But I guess...

But look, this is a hard problem. This goes back to where we started, which is, okay, who has the truth. And it's like, well, how do societies arrive at truth and how do we figure these things out? And our elected leaders play some role in it. We all play some role in it. There have to be some set of public intellectuals at some point that bring rationality and judgment and humility to it. Those people are few and far between. We should probably prize them very highly. Yeah, so celebrate humility in our public leaders. So getting to risk number two, will AI ruin our society? Short version, as you write, if the murder robots don't get us, the hate speech and misinformation will. And the action you recommend in short, don't let the thought police suppress AI. Well, what is this risk of the effect of misinformation in a society that's going to be catalyzed by AI?

Yeah, so this is the social media. This is what you just alluded to. It's the activism kind of thing that's popped up in these companies in the industry. And it's basically, from my perspective, it's basically part two of the war that played out over social media over the last 10 years.

Because you probably remember, social media 10 years ago was basically who even wants this, who wants a photo of what your cat had for breakfast, like this stuff is silly and trivial.

And why can't these nerds figure out how to invent something useful and powerful?

And then certain things happened in the political system. And then it's sort of the polarity on that discussion switched all the way to social media is the worst, most corrosive,

most terrible, most awful technology ever invented. And then it leads to terrible politicians and politics and all this stuff. And that all got catalyzed into this very big kind of angry movement both inside and outside the companies to kind of bring social media to heel. And that got focused in particularly on two topics, so-called hate speech and so-called misinformation. And that's been the saga playing out for the last decade. And I don't even really want to even argue the pros and cons of the sides just to observe that that's been like a huge fight and has had big consequences to how these companies operate. Basically, those same sets of theories that same activist approach, that same energy is being transplanted straight to AI. And you see that already happening. It's why ChatGPT will answer, let's say, certain questions and not others. It's why it gives you the canned speech about whenever it starts with as a large language model, I cannot basically means that somebody has reached in there and told that it can't talk about certain topics. Do you think some of that is good? So it's an interesting question. So a couple of observations. So one is the people who find this the most frustrating are the people who are worried about the murder robots. So in fact, the so-called X-risk people, they started with the term AI safety. The term became AI alignment. When the term became AI alignment is when this switch happened from where worried is going to kill us all to where worried about hate speech and misinformation. The AI X-risk people have now renamed their thing, AI not kill everyone ism, which I have to admit is a catchy term. And they are very frustrated by the fact that the sort of activist driven hate speech misinformation kind of thing is taking over, which is what's happened is taking over the ethics field has been taken over by the hate speech misinformation people. You know, look, would I like to live in a world in which like everybody was nice to each other all the time and nobody ever said anything mean and nobody ever used a bad word and everything was always accurate and honest, like that sounds great. Do I want to live in a world where there's like a centralized thought police working through the tech companies to enforce the view of a small set of elites that they're going to determine what the rest of us think and feel like absolutely not. There could be a middle ground somewhere like Wikipedia type of moderation, there's moderation of Wikipedia, that it's somehow crowdsourced where you don't have centralized elites. But it's also not completely just a free for all because the if you have the entirety of human knowledge at your fingertips, you can do a lot of harm. Like if you have a good assistant that's completely uncensored, they can help you build a bomb. They can help you mess with people's physical well-being, right? Because that information is out there on the internet. So presumably there's, it would be, you could see the positives in censoring some aspects of an AI model when it's helping you commit literal violence. And there's a section, later section of the essay where I talk about bad people doing bad things. And there's a set of things that we should discuss there. What happens in practice is these lines, as you alluded to this already, these lines are not easy to draw. And what I've observed in the social media version of this is, like the way I describe it, is the slippery slope is not a fallacy, it's an inevitability. The minute you have this kind of activist personality that gets in a position to make these decisions, they take it straight to infinity. It goes into the crazy zone almost immediately and never comes back because people become drunk with power. And look, if you're in the position to determine what the entire world thinks and feels and reads and says, like, you're going to take it. And Elon has ventilated this with the Twitter files over the last three months. And it's just crystal clear like how bad it got

there. Now, reason for optimism is what Elon is doing with the community notes. So community notes

is actually a very interesting thing. So what Elon is trying to do with community notes is he's trying to have it where there's only a community note when people who have previously disagreed on many topics agree on this one. Yes, that's what that's what I'm trying to get at is like there's there could be Wikipedia like models or community notes type of models where allows you to essentially either provide context or sensor in a way that's not resist the slippery slope nature. Now, there's another power. There's an entirely different approach here, which is basically, we have AIs that are producing content, we could also have AIs that are consuming content. And so one of the things that your assistant could do for you is help you consume all the content and basically tell you when you're getting played. So for example, I'm going to want the AI that my kid uses to be very child safe. And I'm going to want it to filter for him all kinds of inappropriate stuff that he shouldn't be saying just because he's a kid. And you see what I'm saying is you can implement that. You could do the architecture, you could say you can solve this on the client side. Right. Solving on the server side gives you an opportunity to dictate for the entire world, which I think is where you take the slipper slope to hell. There's another architectural approach, which is to solve this on the client side, which is certainly what I would endorse.

It's AI risk number five, will AI lead to bad people doing bad things? I can just imagine language models used to do so many bad things, but the hope is there that you can have large language models used to then defend against it by more people, by smarter people, by more effective people, skilled people, all that kind of stuff.

Three-point argument on bad people doing bad things. So number one, right, you can use the technology defensively. We should be using AI to build broad-spectrum vaccines and antibiotics for bioweapons, and we should be using AI to hunt terrorists and catch criminals, and we should be doing all kinds of stuff like that. In fact, we should be doing those things even just to go eliminate risk from regular pathogens that aren't constructed by an AI. So there's the whole defensive set of things. Second is we have many laws on the books about the actual bad things. So it is actually illegal to commit crimes, to commit terrorist acts, to build pathogens with the intent to deploy them to kill people.

We actually don't need new laws for the vast majority of these scenarios. We actually already have the laws on the books. The third argument is the minute, and this is sort of the foundational one that gets really tough, but the minute you get into this thing, which you were kind of getting into, which is like, okay, but you need censorship sometimes, right? And don't you need restriction sometimes? It's like, okay, what is the cost of that? And in particular, in the world of open source, right? And so is open source AI going to be allowed or not? If open source AI is not allowed, then what is the regime that's going to be necessarily legally and technically to prevent it from developing, right? And here, again, is where you get into, and people have proposed that these kinds of things, you get into, I would say, pretty extreme territory pretty fast. Do we have a monitor agent on every CPU and GPU that reports back to the government, what we're doing with our computers? Are we seizing GPU clusters to get beyond a certain size? And then, by the way, how are we doing all that globally, right? And if China is developing an LLM beyond the scale that we think is allowable, are we going to invade, right? And you have figures on the AIX risk side who are advocating potentially up to nuclear strikes to prevent this kind of thing. And so here

you get into this thing. And again, you could maybe say this is, you could even say this is what good, bad, or indifferent, or whatever. But like, here's the comparison of nukes. The comparison of nukes is very dangerous, because one is just nukes were just a bomb, although we can come back to nuclear power. But the other thing was like, with nukes, you could control plutonium, right? You could track plutonium. And it was like hard to come by. AI is just math and code, right? And it's in like math textbooks. And it's like there are YouTube videos that teach you how to build it. And like, there's open source, it's already open source, you know, there's a 40 billion parameter model running around already called Falcon Online that anybody can download. And so, okay, you walk down the logic path that says we need to have guardrails on this, and you find yourself in a authoritarian, totalitarian regime of thought control and machine control that would be so brutal that you would have destroyed the society that you're trying to protect. And so I just don't see how that actually works. So yeah, you have to understand my brain is going full steam ahead here, because I agree with basically everything you're saying when I'm trying to play devil's advocate here, because okay, you highlighted the fact that there is a slippery slope to human nature. The moment you sense there's something, you start to sense everything.

That alignment starts out sounding nice, but then you start to align to the beliefs of some select group of people, and then it's just your beliefs. The number of people you're aligning to smaller and smaller as that group becomes more and more powerful. Okay, but that just speaks to the people that censor usually the assholes and the assholes get richer.

I wonder if it's possible to do without that for AI. The one way to ask this question is do you think the base models, the baseline foundation models should be open sourced? Like what Mark Zuckerberg is saying they want to do.

So look, I think it's totally appropriate that companies that are in the business of producing a product or service should be able to have a wide range of policies that they put. And I'll just again, I want a heavily censored model for my eight-year-old. I actually want that. I would pay more money for the ones more heavily censored than the one that's not. And so there are certainly scenarios where companies will make that decision. Look, an interesting thing you brought up is this really a speech issue. One of the things that the big tech companies are dealing with is that content generated from an LLM is not covered under section 230, which is the law that protects internet platform companies from being sued for the user-generated content. And so it's actually, yes, and so there's actually a question. I think there's still a question which is can big American companies actually feel generative AI at all? Or is the liability actually going to just ultimately convince them that they can't do it? Because the minute the thing says something bad, and it doesn't even need to be hate speech, it could just be like an enact, it could hallucinate a product detail on a vacuum cleaner, and all of a sudden the vacuum cleaner company sues for misrepresentation. And there's any symmetry there, right? Because the LLM is going to be producing billions of answers to questions, and it only needs to get a few wrong to have. The loss has to get updated really quick here.

Yeah, and nobody knows what to do with that. So anyway, there are big questions around how companies operate at all. So we talk about those. But then there's this other question of like, okay, the open source, so what about open source? And my answer to your question is kind of like, obviously, yes, the models, there has to be full open source here, because to live in a world in

which that open source is not allowed is a world of draconian speech control, human control, machine

control. I mean, you know, black helicopters with jackbooted thugs coming out, repelling down and seizing your GPU, like, territory. Well, no, no, I'm 100% serious.

That's you're saying slippery slope always leaves that. No, no, no, no, no, no, no, that's what's required to enforce it. Like, how will you enforce a ban on open source?

No, you could add friction to it. Like, hard to get the models, people will always be able to get the models. But it'll be more in the shadows, right? The leading open source model right now is from the UAE. Like, the next time they do that, what do we do? Yeah. Like,

Oh, I see. You're like a 14 year old in Indonesia comes out with a breakthrough.

You know, we talked about most great software comes from a small number of people. Some kid comes out with some big new breakthrough and quantization or something and has some huge breakthrough and like, what are we going to like, invade Indonesia and arrest him?

It seems like in terms of size models and effectiveness of models, the big tech companies will probably lead the way for quite a few years. And the question is of what policies they should use. The kid, the kid in Indonesia should not be regulated, but should Google Meta, Microsoft, OpenAI be regulated? Well, so, but this goes, okay, so when does it become dangerous? Yeah. Right. Is the danger that it's quote as powerful as the current leading commercial model?

Or is it that it is, it is just at some other arbitrary threshold? Yeah. And then by the way, like, look, how do we know? Like what we know today is that you need like a lot of money to like train these things, but there are advances being made every week on training efficiency and, you know, data, all kinds of synthetic, you know, look, I don't even like the synthetic data thing we're talking about, maybe some kid figures are a way to auto generate synthetic data.

That's going to change everything. Yeah, exactly. And so like sitting here today, like the breakthrough just happened, right? You made this point, like the breakthrough just happened. So we don't know what the shape of this technology is going to be. I mean, the big shock, the big shock here is that, you know, whatever number of billions of parameters basically represents at least a very big percentage of human thought, like who would have imagined that? And then there's already work underway. There was just this paper that just came out that basically takes a GPT three scale model and compresses it down to run on a single 32 core CPU. Like who would

have predicted that? Yeah. You know, some of these models now you can run in Raspberry Pis, like today they're very slow, but like, you know, maybe they'll be, you know, perceived, you've really performed, you know, like it's math and code. And here we're back in here, we're back in math and code. It's math and code. It's math, code and data. It's bits.

Mark's just like walking away at this point. Screw it. I don't know what to do with this.

You guys created this whole internet thing. Yeah. Yeah. I'm a huge believer in open source here. So my argument is we're going to have to see, here's my argument is my argument, my full argument is AI is going to be like air, it's going to be everywhere.

Like this is just going to be in text. It already is. It's going to be in textbooks and kids are going to grow up knowing how to do this. And it's just going to be a thing. It's going to be in the air. And you can't like pull this back anymore. You can pull back air. And so you just have to figure out how to live in this world, right? And then that,

and then that's where I think like all this hand wringing about AI risk is basically a complete waste of time because the effort should go into, okay, what are, what are, what is the defensive approach? And so if you're worried about, you know, AI generated pathogens, the right thing

to do is to have a permanent project warp speed, right? And funded lavishly. Let's, let's do a Manhattan, let's talk about Manhattan project for biological defense, right? And let's build AIs and let's have like broad spectrum vaccines where like we're insulated from every pathogen, right? And what, what the interesting thing is, because it's software, a kid in his basement, teenager could build like a system that defends against like the worst, the worst. I mean, and to me, defense is super exciting. It's like, I, if you believe in the good of human nature, that most people want to do good to be the savior of humanity is really exciting.

Yes. Not, okay, that's a dramatic statement, but like to help people, to help people.

Yeah. Okay. What about just to jump around? What about the risk of will AI lead to crippling inequality? You know, because we're kind of saying everybody's life will become better.

Is it possible that the, the rich get richer here?

Yeah. So this goes, it's actually ironically goes back to Marxism. So, because this was the the core claim of Marxism, right? Basically it was that the owner, the owners of capital would basically own the means of production. And then over time, they would basically accumulate all the wealth. The workers would be paying in, you know, and getting nothing in return because they wouldn't be needed anymore, right? Marx is very worried about what he called mechanization or what later became known as automation. And that, you know, the workers would be emissary than the capitalists would end up with, with all. And so this was one of the core core core principles of Marxism. Of course, it turned out to be wrong about every previous wave of technology. The reason it turned out to be wrong about every previous wave of technology is that the way that the self-interested owner of the machines makes the most money is by providing the production capability in the form of products and services to the most people, the most customers as possible, right? The largest, and this is one of those funny things where every CEO knows this intuitively, and yet it's like hard to explain from the outside. The way you make the most money in any business is by selling to the largest market you can possibly get to. The largest market you can possibly get to is everybody on the planet. And so every large company does is everything that it can to drive down prices to be able to get volumes up to be able to get to everybody on the planet. And that happened with everything from electricity. It happened with telephones. It happened with radio. It happened with automobiles. It happened with smartphones. It happened with the PCs. It happened with the internet. It happened with mobile broadband. It's happened, by the way, with Coca-Cola, and it's happened with

like every, you know, basically every industrially produced, you know, good or service, people want you want to drive it to the largest possible market. And then as proof of that, it's already happened, right? Which is the early adopters of like Chet GPD and Bing are not like, you know, Exxon and Boeing. They're, you know, your uncle and your nephew, right? It's just like, it's either freely available online or it's available for 20 bucks a month or something. But you know, these things went, this technology went mass market immediately. And so look, the owners of the means of production, whoever does this, as I mentioned, these trillion dollar questions,

there are people who are going to get really rich doing this, producing these things, but they're going to get really rich by taking this technology to the broadest possible market. So yes, they'll get rich, but they'll get rich having a huge positive impact on.

Yeah, making that making the technology available to everybody.

Yeah.

Right. And again, smartphone, same thing. So there's this amazing kind of twist in business history, which is you cannot spend \$10,000 on a smartphone, right? You can't spend \$100,000. You can't spend like, I would buy the million dollar smartphone, like I'm signed up for it.

Like if it's like, suppose a million dollar smartphone was like much better than a thousand dollar smartphone, like I'm there to buy it, it doesn't exist. Why doesn't it exist? Apple makes so much more money driving the price further down from \$1,000 than they would trying to harvest, right? And so it's just this repeating pattern you see over and over again.

And what's great about it, what's great about it is you do not need to rely on anybody's enlightened right generosity to do this. You just need to rely on capitalist self-interest.

What about AI taking our jobs?

Yeah. So very, very similar thing here. There's sort of a, there's a core fallacy, which again, was was very common in Marxism, which is what's called the lump of labor fallacy.

And this is sort of the fallacy that there is only a fixed amount of work to be done in the world.

And if the, and it's all being done today by people, and then if machines do it,

there's no other work to be done by people. And that's just a completely backwards

view on how the economy develops and grows. Because what happens is not, in fact, that what happens is the introduction of technology into production process causes prices to fall.

As prices fall, consumers have more spending power. As consumers have more spending power, they create new demand. That new demand then causes capital and labor to form into new enterprises

to satisfy new wants and needs. And the result is more jobs at higher wages.

So new wants and needs, the, the worries that the, the creation of new wants and needs at a rapid rate will mean there's a lot of turnover in jobs. So people will lose jobs, just the actual experience of losing a job and having to learn new things and new skills is painful for the individuals.

Two things. One is the new jobs are often much better. So it's actually came up that there was this panic about a decade ago on all the truck drivers are going to lose their jobs, right?

And number one, that didn't happen because we haven't figured out a way to actually finish that yet. But, but the other thing was like, like truck driver, like I grew up in a town that was basically consisted of a truck stop, right? And I like knew a lot of truck drivers and like truck drivers live a decade shorter than everybody else. Like they, it's a, it's a, it's actually like a very dangerous, like they get like literally they have like high racist skin cancer. And on the left side of their, on the left side of their body from, from being in the sun all the time, the vibration of being in the truck is actually very damaging to your, to your physiology. And there's actually a, perhaps partially because of that reason, there's a shortage of people who want to be truck drivers.

Yeah. Like it's not, it's not like, the question always you want to ask somebody like that is, do you want, you know, do you want your kid to be doing this job? And like most of them will

tell you, no, like I want my kid to be sitting in a cubicle somewhere, like where they don't have this, like where they don't die 10 years earlier. And so, so the new jobs, number one, the new jobs are often better, but you don't get the new jobs until you go through the change. And then to your point, the, the training thing, you know, it's always the issue is can, can people adapt? And again, here you need to imagine living in a world in which everybody has the AI assistant capability, right, to be able to pick up new skills much more quickly and be able to have some, you know, be able to have a machine to work with to augment their skills.

It's still going to be painful, but that's the process of life.

It's painful for some people. I mean, there's no, like there's no question it's painful for some people and they're, you know, they're, yes, it's, it's not, again, I'm not a utopian on this and it's not like it's positive for everybody in the moment, but it has been overwhelmingly positive for 300 years. I mean, look, the concern here, the concern, the concern, this concern has played out for literally centuries. And, you know, this is the sort of leadite, you know, the story of the leadites that you may remember, there was a panic in the 2000s around outsourcing was going to take all the jobs. There was a panic in the 2010s that robots were going to take all the jobs. In 2019, before COVID, we had more jobs at higher wages, both in the country and in the world than at any point in human history. And so the overwhelming evidence is that the net gain here is like, just like wildly positive. And most, most people like overwhelmingly come out the other side being huge beneficiaries of this. So you write that the single greatest risk, this is the risk you're most convinced by the single greatest risk of AI is that China wins global AI dominance and we, the United States and the West do not. Can you elaborate? Yeah. So this is the other thing, which is a lot of the sort of AI risk debates today sort of assume that we're the only game in town, right? And so we have the ability to kind of sit in the United States and criticize ourselves and do, you know, have our government like, you know, beat up on our companies and we're figuring out a way to restrict what our companies can do. And, you know, we're going to, you know, we're going to ban this and ban that, restrict this and do that. And then there's this like other like force out there that like doesn't believe we have any power over them whatsoever. And they have no desire to sign up for whatever rules we decided to put in place. And they're going to do whatever it is they're going to do. And we have no control over it at all. And it's China and specifically the Chinese Communist Party. And they have a completely publicized, open, you know, plan for what they're going to do with AI. And it is not what we have in mind. And not only do they have that as a vision and a plan for their society, but they also have it as a vision and plan for the rest of the world. So their plan is what? Surveillance? Yeah, authoritarian control. So authoritarian population control, you know, good old fashioned communist authoritarian control, and surveillance and enforcement, and social credit scores and all the rest of it. And you are going to be monitored and metered within an inch of everything all the time. And it's going to be basically the end of human freedom. And that's their goal. And, you know, they justify it on the basis of that's what leads to peace.

And you're worried that the regulating in the United States will hold progress enough to where the Chinese government would win that race? So their plan, yes, yes. And the reason for that is they, and again, they're very public on this, they have their plan is to proliferate their approach around the world. And they have this program called the digital Silk Road, right, which is building on their Silk Road investment program. And they've got their, they've been

laying, they've been laying networking infrastructure all over the world with their 5G right work with their company, Huawei. And so they've been laying all this fabric, but financial and technological fabric all over the world. And their plan is to roll out their vision of AI on top of that, and to have every other country be running their version. And then if you're a country prone to authoritarianism, you're going to find this to be an incredible way to become more authoritarian. If you're a country, by the way, not prone to authoritarianism, you're going to have the Chinese Communist Party running your infrastructure and having backdoors into it, right, which is also not good.

What's your sense of where they stand in terms of the race towards superintelligence as compared to the United States?

Yeah, so good news is they're behind, but bad news is they, you know, they, let's just say they get access to everything we do. So they're probably a year behind at each point in time, but they get, you know, downloads, I think of basically all of our work on a regular basis through a variety of means. And they are, you know, we'll see they're at least putting out reports. They just put out a report last week of a GPT 3.5 analog. They put out this report, forget what it's called, but they put out this report of this LN they did. And they, you know, the way, when OpenAI puts out, one of the ways they test GPT is they run it through standardized exams like the SAT, right, just how you can kind of gauge how smart it is. And so the Chinese report, they ran their LM through the Chinese equivalent of the SAT, and it includes a section on Marxism and a section on Mao Zedong thought. And it turns out their AI does very well on both of those topics, right? So like this, this alignment thing, communist AI, right? Like literal communist AI, right? And so their vision is like that's the, you know, so, you know, you can just imagine, like you're a school, you know, you're a kid 10 years from now in Argentina or in Germany or in who knows where Indonesia and you ask them, I'd explain to you like how the economy works and it gives you the most cheery upbeat explanation of Chinese style communism you've ever heard, right? So like the stakes here are like really big.

Well, my, as we've been talking about, my hope is not just with the United States, but with just the kitten as basement, the open source LLM. So I don't know if I trust large centralized institutions with super powerful AI, no matter what their ideology, there's a power corrupts. You've been investing in tech companies for about, let's say 20 years and about 15 of which was with Andreessen Horowitz. What interesting trends in tech have you seen over that time? Let's just talk about companies and just the evolution of the tech industry. I mean, the big shift over 20 years has been that tech used to be a tools industry for basically from like 1940 through to about 2010, almost all the big successful companies were picks and shovels companies. So PC database, smartphone, you know, some tool that

somebody else would pick up and use. Since 2010, most of the big wins have been in applications. So a company that starts in an existing industry and goes directly to the customer in that industry. And the early examples there were like Uber and Lyft and Airbnb. And then that model is kind of elaborating out. The AI thing is actually a reversion on that for now, because like most of the AI business right now is actually in cloud provision of AI APIs for other people to build on. But the big thing will probably be an app. Yeah, I think most of the money, I think probably will be in whatever, yeah, your AI financial advisor or your AI doctor or your AI lawyer or take your

pick of whatever the domain is. And what's interesting is the valley kind of does everything. The entrepreneurs kind of elaborate every possible idea. And so there will be a set of companies that like make AI something that can be purchased and used by large law firms. And then there will be other companies that just go direct to market as an AI lawyer. What advice could you give for a startup founder? Just having seen so many successful companies, so many companies that fail also. What advice could you give to a startup founder, someone who wants to build the next super successful startup in the tech space, the Googles, the Apples, the Twitters? Yeah, so the great thing about the really great founders is they don't take any advice. So if you find yourself listening to advice, maybe you shouldn't do it. But that's actually just to elaborate on that. If you could also speak to great founders too. Like what makes a great founder? So what makes a great founder is super smart, coupled with super energetic, coupled with super courageous. I think it's some of those three. Intelligence, passion, and courage. The first two are traits and the third one is a choice, I think. Courage is a choice. Well, because courage is a question of pain, tolerance, right? So how many times are you willing to get punched in the face before you quit? And here's maybe the biggest thing people don't understand about what it's like to be a startup founder is, it gets very romanticized, right? And even when they fail, it still gets romanticized about what a great adventure it was. But the reality of it is most of what happens is people telling you, no, and then they usually follow that with, you're stupid, right? No, I will not come to work for you. I will not leave my cushy job at Google to come work for you. No, I'm not going to buy your products. No, I'm not going to run a story about your company. No, I'm not this, that, the other thing. And so a huge amount of what people have to do is just get used to just getting punched. And the reason people don't understand this is because when you're a founder, you cannot let on that this is happening because it will cause people to think that you're weak and they'll lose faith in you. So you have to pretend that you're having a great time when you're dying inside, right? You're just a misery. But why did they do it? Yeah, that's the thing, it's like it is a level, this is actually one of the conclusions I think is it, I think it's for most of these people on a risk-adjusted basis, it's probably an irrational act. They could probably be more financially successful on average if they just got like a real job in a big company. But there's, some people just have an irrational need to do something new and build something for themselves. And some people just can't tolerate having bosses. Oh, here's a fun thing is how do you reference check founders, right? So you call it, you know, normally you reference check your time hiring somebody as you call the bosses, and you find out if they were good employees. And now you're trying to reference check Steve Jobs, right? And it's like, oh God, he was terrible, you know, he was a terrible employee, he never did what we told him to do. So what's a good reference? Do you want the previous boss to actually say they never did what you told them to do? That might be a good thing. Well, ideally, ideally what you want is, I would like to go to work for that person. He worked for me here, and now I'd like to work for him. Now, unfortunately, most people can't, their egos can't handle that, so they won't say that, but that's the ideal. What advice would you give to those folks in the space of intelligence, passion, and courage? So I think the other big thing is, you see people sometimes who say, I want to start a company, and then they kind of work through the process of coming up with an idea.

And generally, those don't work as well as the case where somebody has the idea first, and then they kind of realize that there's an opportunity to build a company, and then they just turn out to be the right kind of person to do that. When you say idea, do you mean long term big vision, or do you mean specifics of like product? I would say specific, like specifically what, yes, specifics, like what is the, because for the first five years, you don't get to have vision, you just got to build something people want, and you got to figure out a way to sell it to them, right? It's very practical, or you never get to big vision, so.

So the first, you have an idea of a set of products or the first product that can actually make some money. Yeah, like it's got a, first product's got to work, by which I mean like, it has to technically work, but then it has to actually fit into the category in the customer's mind of something that they want. And then, and then by the way, the other part is they have to want to pay for it, like somebody's got to pay the bills, and so you got to figure out how to price it, and whether you can actually extract the money. Yeah. So usually it is much more predictable, success is never predictable, but it's more predictable if you start with a great idea, and then back into starting the company. So this is what we did, you know, we had, most like before we had an escape, the Google guys had the Google search engine working at Stanford, right? The, you know, yeah, actually, there's tons of examples where they, you know, Pierre Omidar had eBay working before he left his previous job.

So I really love that idea of just having a thing, a prototype that actually works, before you even begin to remotely scale. Yeah. By the way, it's also far easier to raise money, right? Like the ideal pitch that we receive is, here's the thing that works, would you like to invest in our company or not? Like that's so much easier than here's 30 slides with a dream, right? And then we have this concept called the idea maze, which our biology student of Boston came up with when he was with us. So, so, so then there's this thing, this goes to mythology, which is, you know, there's a mythology that kind of, you know, these, these ideas, you know, kind of arrive like magic or people kind of stumble into them. It's like eBay with the pest dispensers or something. The reality usually with the big successes is that the founder has been chewing on the problem for five or 10 years before they start the company. And they often worked on it in school, or they even experimented on it when they were a kid. And they've been kind of training up over that period of time to be able to do the thing. So they're like a true domain expert. And it sort of sounds like mom and apple pie, which is, yeah, you want to be a domain expert in what you're doing, but you would, you know, the mythology is so strong of like, oh, I just like had this idea in the shower and now I'm doing it like, it's generally not that. No, because maybe in the shower we had the exact product implementation details. But yeah, usually you're going to be for like years, if not decades, thinking about like everything around that.

Well, we call it the idea maze because the idea maze basically is like, there's all these permutations. Like for any idea, there's like all these different permutations. Who should the customer be? What shape, form should the product have? And how should we take it to market and all these things. And so the really smart founders have thought through all these scenarios by the time they go out to raise money. And they have like detailed answers on every one of those fronts because they put so much thought into it. The sort of the sort of more haphazard founders haven't thought about any of that. And it's

the detailed ones who tend to do much better. So how do you know what to take a leap? If you have a cushy job or happy life? I mean, the best reason is just because you can't tolerate not doing it, right? Like this is the kind of thing where if you have to be advised into doing it, you probably shouldn't do it. And so it's probably the opposite, which is you just have such a burning sense of this has to be done. I have to do this. I have no choice.

What if it's going to lead to a lot of pain? It's going to lead to a lot of pain.

I think that's what if it means losing sort of social relationships and damaging your relationship with loved ones and all that kind of stuff.

Yeah, look, so like it's going to put you in a social tunnel for sure, right? So you're going to like, you know, there's this game you can play on Twitter, which is you can do any whiff of the idea that there's basically any such thing as work life balance and that people should actually work hard and everybody gets mad. But like the truth is, like all the successful founders are working 80 hour weeks and they're working, you know, they've formed very, very strong social bonds with people they work with. They tend to lose a lot of friends on the outside or put those friendships on ice. Like that's just the nature of the thing. You know, for most people, it's worth the trade off. You know, the advantage maybe younger founders have is maybe they have less, you know, maybe they're not, you know, for example, if they're not married yet or don't have kids yet, that's an easier thing to bite off. Can you be an older founder? Yeah, you definitely can. Yeah. Yeah. Many of the most successful founders are second, third, fourth time founders. They're in their 30s, 40s, 50s. The good news of being an older founder is you know more and you know a lot more about what to do, which is very helpful. The problem is, okay, now you've got like a spouse and a family and kids and like you've kind of go to the baseball game and like you can't go to the baseball, you know, and so it's get life is full of difficult choices.

Yes.

Like Andreessen, you've written a blog post on what you've been up to. You wrote this in October 2022, quote, mostly I try to learn a lot. For example, the political events of 2014 to 2016 made clear to me that I didn't understand politics at all. Referencing maybe some of this, this book here. So I deliberately withdrew from political engagement and fundraising and instead read my way back into history and as far to the political left and political right as I could.

So just high level question. What's your approach to learning?

Yeah, so it's basically, I would say it's an auto-divac. So it sort of goes, it's going down the rabbit holes. So it's a combination. So I kind of allude to it in that quote. It's a combination of breadth and depth. And so I tend to, yeah, I tend to, I go broad by the nature of what I do, I go broad, but then I tend to go deep in a rabbit hole for a while, read everything I can and then come out of it. And I might not, I might not revisit that rabbit hole for, you know, another decade.

And in that blog post that I recommend people go check out, you actually list a bunch of different books that you recommend on different topics on the American left and the American right. It's just a lot of really good stuff. The best explanation for the current structure of our society and politics, you give two recommendations, four books on the Spanish Civil War, six books on

deep history of the American right, comprehensive biographies of Adolf Hitler, one of which I read and can recommend, six books on the deep history of the American left. So American right,

American left looking at the history to give you the context.

Biography of Lenin, two of them on the French Revolution. Actually, I have never read a biography on Lenin. Maybe that, that'll be useful. Everything's been so Marx focused.

The Sebastian biography of Lenin is extraordinary.

Victor Sebastian, okay.

It'll blow your mind.

Yeah.

So it's still useful to reach.

It's incredible. Yeah, it's incredible. I actually think it's the single best book on the Soviet Union.

So that, the perspective of Lenin might be the best way to look at the Soviet Union versus Stalin versus Marx versus, very interesting.

So two books on fascism and anti-fascism by the same author, Paul Guthrie.

Brilliant book on the nature of mass movements and collective psychology, the definitive work on intellectual life under totalitarianism, the captive mind, the definitive work on the practical life under totalitarianism. There's a bunch, there's a bunch. And the single best book, first of all, the list here is just incredible. But you say the single best book I have found on who we are and how we got here is the ancient city by Neuma Dennis Faustel de Coulancas. I like it. What did you learn about who we are as a human civilization from that book?

Yeah. So this is a fascinating book. This one's free. It's free, by the way. It's a book from the 1860s. You can download it or you can buy prints of it. But it was this guy who was a professor at the Sorbonne in the 1860s. And he was apparently a savant on Greek and Roman antiquity.

And the reason I say that is because his sources are 100% original Greek and Roman sources.

So he wrote basically a history of Western civilization from on the order of 4,000 years ago to basically the present times entirely working on original Greek and Roman sources.

And what he was specifically trying to do was he was trying to reconstruct from the stories of the Greeks and the Romans. He was trying to reconstruct what life in the West was like before the Greeks and the Romans, which was in the civilization known as the Indo-Europeans. And the short answer, and this is sort of circa 2000 BC to 500 BC, kind of that 1500 year stretch where civilization developed. And his conclusion was basically cults. They were basically cults. And civilization was organized into cults. And the intensity of the cults was like a million fold beyond anything that we would recognize today. Like it was a level of all-encompassing belief and an action around religion that was at a level of extremeness that we wouldn't even recognize it. And so specifically, he tells the story of basically there were three levels of cults. There was the family cult, the tribal cult, and then the city cult as society scaled up. And then each cult was a joint cult of family gods, which were ancestor gods, and then nature gods. And then your bonding into a family, a tribe, or a city was based on your adherence to that religion. People who were not of your family tribe city worshiped different gods, which gave you not just the right with the responsibility to kill them on site. Right. So they were serious about their cults. Hard core. By the way, shocking development, I did not realize there's a zero concept of individual rights. Even up through the Greeks, and even in the Romans, they didn't have the concept of individual rights. The idea that as an individual, you have some rights, just like noop. And you look back and you're just like,

wow, that's just like crazily fascist in a degree that we wouldn't recognize today. But it's like, well, they were living under extreme pressure for survival. And the theory goes, you could not have people running around making claims to individual rights when you're just trying to get like your tribe through the winter, right? Like you need like hardcore command and control. And actually, through modern political lens, those cults were basically both fascist and communist. They were fascist in terms of social control, and then they were communist in terms of economics. But you think that's fundamentally that like pull towards cults is within us? My conclusion from this book, so the way we naturally think about the world we live in today, is like we basically have such an improved version of everything that came before us, right? Like we have basically, we figured out all these things around morality and ethics and democracy and all these things. And like they were basically stupid and retrograde and were like smart and sophisticated. And we've improved all this. I after reading that book, I now believe in many ways the opposite, which is no, actually, we are still running in that original model. We're just running in an incredibly diluted version of it. So we're still running basically in cults. It's just our cults are at like a thousandth or a millionth the level of intensity, right? And so just to take religions, the modern experience of a Christian in our time, even somebody who considers him a devout Christian is just a shadow of the level of intensity of somebody who belonged to a religion back in that period. And then by the way, it goes back to our discussion, we then sort of endlessly create new cults. Like we're trying to fill the void, right? And the void is a void of bonding. Okay. Living in their era, like everybody living today, transporting that era would view it as just like completely intolerable in terms of like the loss of freedom and the level of basically fascist control. However, every single person in that era, and he really stresses this, they knew exactly where they stood. They knew exactly where they belonged. They knew exactly what their purpose was. They know exactly what they needed to do every day. They know exactly why they were doing it. They had total certainty about their place in the universe. So the question of meaning, the question of purpose was very distinctly clearly defined for them. Absolutely. Overwhelmingly, undisputably, undeniably. As we turn the volume down on the cultism, we start to, the search for meaning starts getting harder and harder. Yes, because we don't have that. We are ungrounded. We are uncentered, and we all feel it, right? And that's why we reach for, you know, it's why we still reach for religion. It's why we reach for, you know, people start to take on, you know, let's say, you know, a faith in science, maybe beyond where they should put it. You know, and by the way, like sports teams are like a, you know, they're like a tiny little version of a cult and, you know, apple keynotes are a tiny little version of a cult, right? You know, political, you know, and there's cult, you know, there's full blown cults on both sides of the political spectrum right now, right? You know, operating in plain sight. But still not full blown compared as to what it was. Compared to what it used to be. I mean, we would today consider full blown, but like, yes, they're at like, I don't know, a hundred thousandth or something of the intensity of what people had back then. So we live in a world today that in many ways is more advanced and moral and so forth. And it's certainly a lot nicer, much nicer world to live in, but we live in a world that's like very washed out. It's like everything has become very colorless and gray as compared to how people used to experience things, which is I think why we're so prone to reach for drama. There's something in us deeply evolved where we want that back.

And I wonder where it's all headed as we turn the volume down more and more. What advice would you give to young folks today? In high school and college, how to be successful in their career, how to be successful in their life? Yeah. So the tools that are available today are just like I sometimes bore kids by describing what it was like to go look up a book, to try to discover a fact in the old days, the 1970s, 1980s, go to the library and the card catalog and the whole thing. You go through all that work and then the book is checked out and you have to wait two weeks. To be in a world not only where you can get the answer to any question, but also the world now, the AI world where you've got the assistant that will help you do anything, help you teach, learn anything. Your ability both to learn and also to produce is just like, I don't know, a million fold beyond what it used to be. I have a blog post I've been wanting to write. It was I call out, where are the hyperproductive people? That's a good question. With these tools, there should be authors that are writing like hundreds or thousands of like outstanding books. Well, with the authors, there's a consumption question too. But yeah, well, maybe not. Maybe not. You're right. But so the tools are much more powerful.

Artists, musicians. Why aren't musicians producing a thousand times the number of songs? The tools are spectacular.

What's the explanation and by way of advice? Is motivation starting to be turned down a little bit or what? I think it might be distraction. It's so easy to just sit and consume that I think people get distracted from production. But if you wanted to, you know, as a young person, if you wanted to really stand out, you could get on like a hyperproductivity curve very early on. There's a great story in Roman history of plenty of the elder who was this legendary statesman. He died in the Vesuvius eruption trying to rescue his friends. But he was famous both for basically being a polymath, but also being an author. And he wrote apparently like hundreds of books. Most of which have been lost, but he like wrote all these encyclopedias. And he literally like would be reading and writing all day long, no matter what else is going on. And so he would like travel with like four slaves and two of them were responsible for reading to him and two of them were responsible for taking dictation. And so like he'd be going across country and like literally he would be writing books like all the time. And apparently they were spectacular. There's only a few that have survived, but apparently they were amazing. So there's a lot of value to being somebody who finds focus in this life.

Yeah. And there are examples like there are, you know, there's this guy, Judge, what's his name, Posner, Posner, who wrote like 40 books and was also a great federal judge. You know, there's our friend Balgi, I think it's like this. He's one of these, you know, where his output is just prodigious. And so it's like, yeah, I mean, with these tools, why not? And I kind of think we're at this interesting kind of freeze frame moment where like this, these tools are not everybody's hands and everybody's just kind of staring at them trying to figure out what to do. Yeah. The new tools. We have discovered fire.

Yeah. And trying to figure out how to use it to cook. Yeah.

Right. You told Tim Ferriss that the perfect day is caffeine for 10 hours and alcohol for four hours. You didn't think I'd be mentioning this, did you? It balances everything out perfectly, as you said. So perfect. So let me ask, what's the secret to balance and maybe two happiness in life? I don't believe in balance. So I'm the wrong person to ask. Can you elaborate why you don't believe in balance? I mean, maybe it's just, and I look, I think people,

I think people are wired differently. So I think it's hard to generalize this kind of thing, but I'm much happier and more satisfied when I'm fully committed to something. So I'm very much in favor of all of imbalance. Yeah.

In balance. And that applies to work, to life, to everything.

Yeah. Now, I happen to have whatever twist of personality traits lead that in non-destructive dimensions, including the fact that I've actually, I now no longer do the 10-4 plan. I stopped drinking. I do the caffeine, but not the alcohol. So there's something in my personality whatever maladaptation I have is inclining me towards productive things, not unproductive things. So you're one of the wealthiest people in the world. What's the relationship between wealth and happiness? Oh, money and happiness. So I think happiness, I don't think happiness is the thing to strive for. I think satisfaction is the thing.

That's, that just sounds like happiness, but turned down a bit.

Happiness is a walk in the woods at sunset, an ice cream cone, a kiss. The first ice cream cone is great. This 1000th ice cream cone, not so much. At some point, the walks in the woods get boring. What's the distinction between happiness and satisfaction?

Satisfaction is a deeper thing, which is like having found a purpose and fulfilling it, being useful. So just something that permeates all your days, just this general contentment of being useful. Then I'm fully satisfying my faculties, then I'm fully delivering, right, on the gifts that I've been given, that I'm making the world better, that I'm contributing to the people around me, right? And then I can look back and say, wow, that was hard, but it was worth it. I think generally it seems to lead people in a better state than pursuit of pleasure, pursuit of quote unquote happiness. Does money have anything to do with that?

I think the founders and the founding fathers in the US threw this off kilter when they used the phrase pursuit of happiness. I think they should have said, they said pursuit of satisfaction, we might live in a better world today. Well, you know, they could have elaborated on a lot of things. They could have tweaked the second amendment. I think they were smarter than realized. They said, you know what, we're going to make it ambiguous and let these, these humans figure out the rest, these tribal cult like humans figure out the rest.

But money empowers that.

So I think, and I think there, I mean, look, I think Elon is, I don't think I'm even a great example, but I think Elon would be the great example of this, which is like, you know, look, he's a guy who from every, every day of his life, from the day he started making money at all, he just plows into the, into the next thing. And so I think, I think money is definitely an enabler for satisfaction. Let's, let's say money applied to happiness leads people down very dark paths. Very destructive avenues. Money applied to satisfaction, I think could be, it is a real tool. I always let, by the way, I was like, you know, Elon is the case study for behavior. But the other thing that so he's really made me think is Larry, Larry Page was asked one time what his approach to philanthropy was. And he said, oh, I'm just my philanthropic plan is just give all the money to Elon. Right? Well, let me actually ask you about Elon. What, what are your, you've interacted with quite a lot of successful engineers and business people. What do you think is special about Elon? We talked about Steve Jobs.

What, what do you think is special about him as a leader as an innovator?

Yeah. So the, the core of it is he's a, he's, he's back to the future. So he is, he is doing

the most leading edge things in the world, but with a, with a really deeply old school approach. And so to find comparisons to Elon, you need to go to like Henry Ford and Thomas Watson and Howard

Hughes and Andrew Carnegie, right? Leland Stanford, Johnny Rockefeller, right? You need to go to the,

what we're called the bourgeois capitalists, like the hardcore business owner operators who basically built, you know, basically built industrialized society, vendor built. And it's a level of hands-on commitment and depth in the business, coupled with an absolute priority towards truth and towards kind of put science and technology down to first principles that is just like absolute. It was just like unbelievably absolute. He really is ideal that he's only ever talking to engineers. Like he does not tolerate, he has the most bullshit talents anybody have ever met. He wants ground truth on every single topic. And he runs his businesses directly day to day devoted to getting to ground truth in every single topic. So you think it was a good decision for him to buy Twitter? I have developed a view in life did not second guess Elon Musk. I know this is going to sound crazy and unfounded, but well, I mean, he's got a quite a track record. I mean, look, the car was a crazy, I mean, the car was, I mean, look, he's done a lot of things that seem crazy, starting a new car company in the United States of America.

The last time somebody really tried to do that was the 1950s, and it was called Tucker Automotive, and it was such a disaster. They made a movie about what a disaster it was. And then Rockets, like who does that? Like that's, there's obviously no way to start in a rocket company like those days are over. And then to do those at the same time. So after he pulled those two off, like, okay, fine. Like, this is one of my areas of like, whatever opinions I had about that, that is just like, okay, clearly or not relevant. Like this is you just, at some point, you just like bet on the person. And in general, I wish more people would lean on celebrating and supporting versus deriding and destroying. Oh, yeah. I mean, look, he drives resentment, like it's like he is a magnet for resentment. Like his critics are the most miserable, like resentful people in the world. Like it's almost a perfect match of like the most idealized, you know, technologists, you know, of the century coupled with like just his critics are just bitter as can be. I mean, it's sort of very darkly comic to watch.

Well, he fuels the fire of that by being an asshole on Twitter at times. And which is fascinating to watch the drama of human civilization, given our cult roots, just fully on fire.

He's running a cult. You could say that very successfully. So now, now that our cults have gone, and we search for meaning, what do you think is the meaning of this whole thing? What's the meaning of life, Mark Andreessen? I don't know the answer to that. I think the meaning of the closest I get to it is what I said about satisfaction. So it's basically like, okay, we were given what we have, like we should basically do our best. What's the role of love in that mix? I mean, like, what's the point of life if you're, yeah, without love, like, yeah. So love is a big part of that satisfaction. Look, like taking care of people is like a wonderful thing. Like, you know, a mentality, you know, there are pathological forms of taking care of people, but there's also a very fundamental, you know, kind of aspect of taking care of people.

Like, for example, I happen to be somebody who believes that capitalism and taking care of

people are actually, they're actually the same thing. Somebody once said, capitalism is how you take care of people you don't know, right? Right. And so like, yeah, I think it's like deeply woven into the whole thing. You know, there's a long conversation to be had about that, but yeah. Yeah, creating products that are used by millions of people and bring them joy in smaller big ways. And then capitalism kind of enables that. Encourages that.

David Friedman says there's only three ways to get somebody to do something for somebody else.

Love, money, and force. Love and money are better. Yeah. That's a good ordering, I think.

We should bet on those. Try love first. If that doesn't work, the money. And then force, well,

don't even try that one. Mark, you're an incredible person. I've been a huge fan. I'm glad

to finally get a chance to talk. I'm a fan of everything you do, everything you do, including

on Twitter. It's a huge honor to meet you, to talk with you. Thanks again for doing this.

Awesome. Thank you, Lex. Thanks for listening to this conversation with Mark Andreessen.

To support this podcast, please check out our sponsors in the description. And now let me

leave you with some words from Mark Andreessen himself. The world is a very malleable place.

If you know what you want and you go for it with maximum energy and drive and passion,

the world will often reconfigure itself around you much more quickly and easily than you would think.

Thank you for listening and hope to see you next time.

You