

(no title)

41 MIN · YOUTUBE · [HTTPS://YOUTU.BE/F_7M4HC-USM?SI=DUEAVNNHARDI5N-F](https://youtu.be/F_7M4Hc-usM?si=DUEAVNNHARDI5N-f)
https://youtu.be/F_7M4Hc-usM?si=DUEAVNNhArdI5n-f

SUMMARY

Sam Olen discusses the evolution of startups and AI, emphasizing the significant changes in how startups can leverage AI technology compared to previous years. He reflects on his experiences with OpenAI and the transformative potential of scaling AI models, while also addressing the challenges of human organization and the future of education in a world increasingly influenced by AI.

- The startup landscape has drastically changed due to advancements in AI, allowing for ambitious projects that were previously impossible.*
- Olen believes that the best startup ideas today are likely not obvious and may be discovered by others rather than assigned as tasks.*
- Scaling has emergent properties that often yield unexpected benefits, which many in the field underestimate.*
- The importance of clear goals and plans is crucial when organizing human teams, especially in rapidly evolving tech environments.*
- Olen reflects on the need for education systems to adapt to the realities of AI, warning against teaching methods that do not foster critical thinking.*
- He discusses the potential for AI to become a utility, similar to electricity, and the importance of democratizing access to AI technology.*
- The conversation touches on the future of economic models, suggesting that ownership stakes in companies may be more beneficial than universal basic income.*
- Olen highlights the ongoing compute shortage and its implications for the future of AI development and demand.*

Please join me in welcoming Sam Olen. [applause] This class was designed as an inspiration from a you know from a set of different experiences uh while I was a student here. One of them was Terry Wintergrad's uh intro seminar CS47N computers and the open society. Uh but a second one that was a pretty formative experience uh for me and a lot of my friends and peers on campus at the time in 2014 was uh CS183 how to start a startup by SAM. Um and so it's really cool to have you back. Uh what's it like? How how's it feeling for you to be back? I was thinking as I was walking in, if I had just a little more time, I would do uh an update to that class because I think everything about starting a startup has changed so much and I have not seen anyone do a good version of how you're supposed to make a startup now. Uh so I had that like just walking in here I had that like ah it'd be fun to do it again.

>> So uh timeline wise yeah you you taught that in 14 I think open was founded in 2015 is that right? >> 16 basically 16. Okay. So, so then you went, you know, it was like you were it it it felt to me from the from an observer perspective that you had like come up with your working theory for how to do it right and then you went and tried to implement it. Is that is that a spare assessment or is that not the case? Obi was like the

strangest startup of the last maybe couple of decades in the Silicon Valley because it started as a research lab. It was it was really not a company at all, >> right? Um, and that the kind of normal course of of startups is that you start a product company and then it like grows for a while and then growth slows down and then you start a research lab and you like bolt that on and you try to figure out the next thing to do. And we were the opposite of that.

We were a research lab first that later had to bolt on a startup, >> right? >> And uh I don't really recommend that. It's kind of an unusual thing, but that that's not quite what I meant. What I meant is like we still followed the preAI rules of a startup because we were trying to make AI. We didn't have it yet. >> But now like watching what the best startups do is so different than how startups worked even a couple of years ago. Um that I think someone I'm probably not going to do it. Someone should do that class again. And what would be the biggest updates you you'd make based on your data? Um, you with like an affordable amount of spend on tokens, you can do what a 100 person incredibly great engineering team would do as a startup and that was just totally impossible. That was like not in the set of options for a startup and now it is.

So, so I think what you can take on, uh, the level of ambition you can have, the speed of which you can move, the amount of stuff you can do at once, uh, is just totally different. And, um, does that change the shape of the problems you feel like you'd assign at the end of the class for people to attack, you know, at the end of that quarter if you were teaching it again? I don't think assigning problems to attack ever works because if you like if I can think of a problem, if I can think of like a really great startup idea, uh if it's like obvious enough to me, uh then it's probably obvious to a lot of people.

When we started OpenAI, we were we were like the uh you know, one of maybe generously speaking four AGI efforts in the world, right? And you want to find something like that. And I'm sure that there exists something today that just wasn't possible at all pre like automated coding era uh that is totally unobvious that will be you know a multi- trillion dollar market soon uh and that only four companies are working on right now. But I don't know what that is. It's much more likely you all know what that is than I know what that is. I just you know my brain is like taken over by open AAI. Um but you know the kind of idea someone can assign you to work on is probably not what you want.

>> Yep. Um, okay. So, that that's fair. Um, but I think it would be helpful since this is a systems class to maybe uh reason about a particular problem that you have to reason through so that they can then apply the shape of the techniques used to break down from a systems perspective that problem into solutions to their own problem. >> Yeah. Um and a and a concept that uh you had started to tease in the class you know back in 2014 and then uh clearly you've talked about publicly over the years is um scale right scale is its own beast it's it's you know quantity is its own quality what scale as a concept has been something it seems like you've um empirically investigated in all kinds of ways over the last 10 years.

Um could you help help us first unpack like what you mean by scale now 10 years later how would you deconstruct that as a systems design uh attribute to apply whether it's a as as a tool um can can we start there yes uh so I don't know why the following observation is true I offer no theory that I find satisfying to explain it and

that makes me a little bit nervous to suggest trust you follow it, but I'm going to anyway because empirically it does seem to be true, which is all of the most interesting things I have observed in my career in watching other uh things happen. All of the most interesting ones uh have had something to do with emergent properties that scale or scale continuing to provide returns far beyond what the consensus thinks will work. And this obviously happens with like scaling loss for AI models. Um but this happens with uh you know getting more smart people together to think about one problem. This h in a in a research setting. Um this happens with uh companies and the sort of economy of scale. You can get all the in all these different ways. I really learned this at Y Combinator when uh it became clear to me that everybody was saying, "Oh, Y Combinator's gotten too big. It should shrink. We should fund less companies per batch." you know, the best times of Y Combinator when it was like 10 companies per batch. And a lot of like very smart people were saying this and and it was like tempting because it would have been like much less work. And the theory was that, you know, the best companies are always kind of obvious and then you fund the rest and it's not as helpful. Um, but a huge part of the magic of what made YC work were uh was the sort of the network effects inside of the batch and that was an emergent property at scale that just hadn't been discovered before. No one had tried to fund startups at scale in the same way and and thus no one had ever happened upon this observation of when you do that um there's there's something important that happens that just didn't exist at all at the 110th to 1/100th of a scale.

There's a bunch of other examples like this. Uh I and I'll skip them in the interest of time, but I I would say again I offer no explanation for why, but empirically speaking, when you find a time that you can push on, you can push something to a scale people have not tried before and it's already working in some interesting way at the smaller scale. More often than not, that seems to be a good idea. And it also seems to be something that most people don't do enough.

>> And I don't offer an explanation for this either, but like in, you know, when we were like, we're really going to scale AI models. Um, all of the like geniuses in the field, most of them were, oh, this isn't really working. You know, that's that's barely a scientific result. It's not interesting that it gets better at scale. You've already shown that. Why keep scaling it? I mentioned the YC example. Um, I've seen a lot of startup founders where they're like, well, you know, there might be something interesting that would happen if I scaled this up, but I'm a little worried about it for non-specific reasons. And again, looking back at like a huge data set of people that have scaled their companies in all these different ways.

There's almost always interesting stuff there. So, I think directionally that's like an interesting thing to push on and severely underexplored. Um on the systems design part of that uh I think one reason people don't do it as much is stuff breaks uh at an accelerating rate and in an unpredictable way as you scale it and if you are going to really scale something um it's always like a little bit broken.

there are always like very smart people who say why you shouldn't do this you know don't get too ambitious don't get too big let's try this smaller and so breaking that down as a systems problem I use the thing of when we were like scaling up AI models there was technically can we do this at all this seems crazy like no one had ever thought about trying to do a run across 10,000 or 100,000 GPUs and that was going to require stacks of engineering talent um there was the capital requirements and what was going to take to do this and like how is

there ever going to be a business how can you think about taking this risk uh there was this sort of like cultural stuff of researchers saying well if we're going to get all this comput something why not have to divide it up among all these all these projects and this also happens in kind of every area I've looked at almost every area for scale and breaking it down into the sort of each difficult area or each reason not to do it and trying to address them one at a time that's been really important. Um, I'm going to push on that a little bit because there's very few people who've been able to sort of repeatedly scale new products and systems the way uh the OpenAI team has over the years. But it seems like one of the issues is there are all these prior conditioning sort of mental models and expectations humans have. And you said things break. And one of the things it seems often breaks that's hard the hardest to refactor is is human the human side of the the systems design, right? Wherever there's human implementers or there's uh human participants in that. And so what have you learned about humans at scale like organizing humans at scale to participate in a system that may not be uh like just a redo of some past system that they they get naively on at a priority on first blush. Um, I think like clear a clear goal, a clear plan to get there. Uh, and like a clear answer to the way that you're going to get there and kind of how you're going to make decisions along the way. That's that's very important. So, um, you know, if we go back to the example of when we decided to scale up models, there were a lot of people who were like, ah, this isn't really going to work. It's going to have these problems. It's also not, you know, we need a more diversified portfolio. But once we say no, we're going to make a bet on scaling deep learning, like that's our thing. If we're wrong, we'll fail, but we're going to do that. Here's why we're going to do that. Here's what we believe about what the state of the world could be like if we get there. Uh, that's very powerful.

And then for whatever reason, um, we did not evolve to be good at thinking about exponentials. People have a hard time imagining that scaling laws are going to continue exponentially, that revenue will grow exponentially, that an organization can take on exponential complexity. And in my experience, it takes a lot of time to really reason through first principles with people about why why that can happen. Can we take two examples uh to walk through that? The first being tach and the second being codeex. You know, both of these have transformed. Can everyone hear? I'm going to try to project it.

Yeah. Okay. Um so let me put in a frame and you can challenge both the assumption and then we can hopefully reason to example what happened. In the case of chat GP you know for a long time in scaling of models a big mental block that seem to be prevalent in the space is what are these things going to be useful for this is you know it's a research uh sort of solution solution chasing a problem research first approach. It's not a product. Um and then you know chat GPT came out and it proved to the world that you know that chat experience was a killer app for general models um at scale for consumers and then a couple of years later you know it's clear that coding has been the killer enterprise app. So what how would you compare and contrast the systems you guys used to discover those use cases ship them scale them monetize them any any salient learnings from those two systems? Yes. Um, so we had made GPT3 and we needed to make money cuz we wanted to go scale up to, you know, a billion and multi-billion dollar computers and we had GPT3 and it was kind of interesting.

It was a cool demo but we couldn't figure out a product to build around it and we had been thinking thinking we just couldn't do it. We had tried a few things. They they hadn't worked. Um, and so we knew the models were gonna get better, but we also wanted to like start a revenue engine sooner. And we said, well, since we can't

figure out what product to build, we're just going to put this into an API and we're going to hope that somebody else can figure out what product to build. And so we launched in like, I don't know, something in the summer of 2020 the GPT3 API. And initially, it kind of got no traction at all. And then about a month later, randomly, as far as we can tell, it went viral on Twitter on the same day, uh, a few different developers kind of found got it to do something cool, posted it, other people started trying, and then like a lot of people started trying the API. Um, but it was shockingly bad. If you go back and use GPT3 or 3.5 um you will be astonished at how bad the models were then uh relative to the amount of excitement they generated at the time. Uh so people tried all these things and really the only business that people got to work in a significant way with GPT3 was copyrighting. Um and that was like not that great and not that exciting and we were kind of like you know h it's just going to have to wait for a better model. But although no that was the only business that was working, developers had figured out how to like put in a prompt and get and be able to chat with it. And we saw this a lot like more people were using they couldn't get the API to work for their business, but they were using their API key to just chat.

And we said, well, we can build a good chatbot. People clearly want that. And we had a new model. We actually had GPT4 done, but we had a new model we were ready to release in between called 3.5. And we had figured out a new kind of post training where we could get the models to do like a good job with instruction following so it can make it easier to chat with. And we said, well, you know, the API is not working great.

Maybe it was like a 10 or a \$20 million run rate kind of business, but there is this thing that people love. Uh, and under the YC principle of see what your users love and do that, we said we'll we'll build a chatbot around it. And we put that out and we still didn't think it was going to do that well. Uh there was it was really meant as like a research demo uh to convince other people that they should build chat light products and pay us for the API, but that went like crazy viral. And another thing I had learned from YC is when something really starts growing and it's not very good, you have like a guaranteed hit on your hands. And so we had like five days where the traffic would shoot up, fall off, and everybody be like, "Well, that was just a hype cycle." But then the next day it would get to a higher peak, fall off again later in the day. People would say that's a hype hype cycle. By the fourth or fifth day, I was like, I know how this works. I know what's going to happen. Like, we have the potential here >> at a killer product. Um, and we knew we could make it much better. We knew we could we knew we had GPT4. We knew we could keep scaling. Um, but by that fifth day, we got everybody together and said, "This is an emergency. This is a good kind of emergency, but we have to build a company and a product all at once.

Uh we then had like two months of crazy scaling. Uh and then we said, you know, we have to figure out a business model later. For now, we're just going to charge people so that we don't like run out our compute bills. But that's obviously not the long-term answer. That also turned out just to work. Um and that was the story of Chinch. And then there was so much utility that people just had not gotten over the activation energy to find that that has worked really well. Um and then codeex.

Actually the plan before chatbt was that we were going to go all in on code. >> Um we knew these models could write code. Uh we knew that they could be really and we knew that that would be like a valuable area. But then

we had this incredibly exciting thing happen. Um but our kind of internal belief at the time was that coding was how these models would control things on computers and robots were how these models would control things in the physical world.

And if you made a smart enough model that had sort of the actuators of writing code and robot and driving a robot, you could then kind of actually get this intelligence to do stuff for you in the world. >> So, uh, then it took us a while to get there. And then I think codeex got really good by early this year, but with 5.5 is when we saw this real inflection point where people are now like doing just incredible things with it. And um you know that we earlier in the class we've talked about how the capabilities pipeline uh is starting to look is starting to become somewhat more legibly standard across different research groups. You got you know pre-training mid-training post training. Then you got the RL and supervised feedback loop. Is do you think that's roughly like the shape of the pipeline that allowed codeex to you know go through a capability jump and that will basically stay stable now and consistent or are we going to go through a major rewrite of that pipeline? I think that is definitely the current pipeline. I expect we will go through a major rewrite. I don't know when it'll happen or exactly how. Um, but it is a little odd to me that it's so happens as a pipeline and doesn't quite feel like the optimal solution. Um, what would be an optimal solution in your head?

>> I think that's a research problem for the AIS to figure out. Um, I think we're at a point where and we've set this goal that by September of this year, we will use 500,000 A100 equivalent GPUs, like a lot of computing power, as an AI research intern, and by March of 2028 that we will have a full end to end very talented researcher like figuring out complete new architectures. Um, so I think we are going to get like with the current pipeline, the current architectures, I think we're going to get over the line of when AIs can do incredible incredible work. Um, you know, one of the things that you you just described there is you you we we've been talking a lot in the class about systems frameworks and analogies to make concepts from one domain legible to other people who may not have all the context in another and that sometimes because of the translation problem, you know, reasoning by analogy is not helpful because then errors compound. Yeah. Um right there you said you know our goal is to try to use it as an AI intern which obviously is a very useful metaphor within the context of you know Silicon Valley a class that understands how these pipelines work and so on and then as as you scale actually that metaphor globally people who might not have all that context go start analogizing these models in ways that they shouldn't be like how should we think about the limits of of that of of what are the limits to scale of um what are the product analogies the research analogies you find most useful within the valley and which one of the what have you found about the limits of those analogies scaling and now how do you navigate between those two problems?

I I've been very interested in studying how like I think what is happening is we are we are in the process of creating a new utility. This doesn't happen very often. you know, electricity is utility, internet's a utility, there water, I guess there's not a lot of these. Uh, and so there are not a lot of examples that we can study for good metaphors or learnings about how to explain this to the world. Um, but I was recently looking at what happened when electricity became a utility. And it's a good analogy for many reasons. It's imperfect, of course, too. But the electricity companies, at least the ones I could find information about, they didn't talk about selling electricity cuz no one knew what that was or why they wanted. It sounds like very scary.

It's this thing that's like going to come into your house and it could kill you in this like gruesome way and you know it feels sort of like very different than the world before. Uh and maybe they tried to sell electricity or market electricity at first. I don't know. But in any case, that didn't work. And then what they started marketing selling to people was light at night. You know, we are going to what you are getting from us is not electricity. It's light at night. By the way, you can use the same thing that lets you get light for all these other things. But people are like, well, why would I want that? And they're like, well, you know, it'll wash your clothes for you someday. And no, no, it won't. I can't. That's too far of a jump for me, >> right?

>> Um, so I don't know what our analogy for this should be. Um, but I suspect that even if even if we're totally right and intelligence is going to become this new utility that every company, every customer, every government just needs access to and is going to use in all sorts of incredible ways and you will have like a OpenAI token subscription that you will plug into everything and use to access everything and you have running for you all the time and doing this amazing stuff. I kind of don't think at least right now the right way for us to analogize that is we're selling intelligence because people are just like somehow not resonating. I don't know what our equivalent of we're selling you light at night is going to be. But I think if we're going to become a new utility, we need to find a way to explain to the world what it means to have this like intelligence pike that you can just do whatever you'd like with >> it. So um one question that has emerged an emerging property of this class of of having a diversity of different speakers is that the utility analogy has come up several times but in reference to different things. So Jensen likened like compute to a utility um and why there should be access and so on and talked about how Stanford should pull budget and so on and and and procure that as a utility for everybody on campus whereas you just likened the intelligence part to util are both of these things true is one of them true one is one more likely to be true how should people reason about compute as a utility versus tokens as a utility and and by comput I mean here chips versus tokens does that make sense >> I think as a consumer as like a business or an individual um you will think in something closer to tokens or probably even one level up from tokens. I don't think you'll care very much about you know where the hardware is, what particular chip it is, what's powering it. I think that stuff will be abstracted out and what you will care about is when you're interacting with the system. Um can you use it a lot? Is it cheap? Is it doing a good job? Um so right now it's like tokens. It may get as we move into a world where we all just have like this constant agent running for us, being useful to us all of the time. Um, you may think about it as even one level up.

But yeah, my my guess is is you when you like pay for your cell phone bill, you're like, "All right, I'm buying access to airtime and some number of gigabytes and, you know, it's going to do all these things and I'll use all these apps and whatever else." But like what you think about paying for the kind of internet utility in this case is just like access to the whole system and the particular hardware at the base station and how it connects to the internet. You don't think about that as much.

>> Um I know I could nerd out about utility infrastructure for a long time but I want to make sure we switch a little bit to being relevant for the students. Usually we have uh questions where we're not hearing those today unless you're comfortable. Oh, okay. Great. How about that? Improv. Okay. Uh so one final question to start getting the creative juices flowing is um the final project for this class or fiber 183 is the oneperson frontier lab. So everybody here is working on projects where they're simulating being an individual uh as a lab with access to all the right tools. They've got hundreds of thousands of dollars of credits from Cloudflare. I think we've got

some open AI tokens maybe. But there's a bunch of compute at their disposal. Um, what would you, if you were in the class, what would you be working on for your oneperson Frontier Lab project? First of all, I think that's an awesome project. Um, I think this is top of mind because uh you we we were just like talking about utility frame frameworks. I think there's a lot of very smart people working on uh great training ideas and we're going to have incredible models.

No matter what you all do, we're going to have incredible models. I promise here uh like pretty quickly but I I think we have not invested enough in being able to deliver at scale huge amounts of cheap intelligence. So maybe I would go work on like the inference part of the stack >> and how are we going to get this incredible intelligence to be cheap and abundant? Uh I think that's underinvested in and and I think all of the frontier labs are going to have to become inference companies to a significant degree. Um, okay.

It might be too late to pivot your projects, but better late than never. >> Work on whatever you want to work on. [laughter] >> Uh, okay. Let's start taking questions and I'm going to moderate and try to be not, you know, please try to be productive and not spicy, etc. Remember, it's a CS class, but up to you Sam is fine. >> Oh, we've got questions. Oh, perfect. All right. First one, the questions about your views on Yan Lun's view that LLMs are a dead end. Um, first of all, in terms of achieving human level intelligence, these models have already far surpassed human intelligence in some ways and then they're wildly worse than others. Like for example, they seem much worse than people are at very long horizon kind of high judgment signal and tasks.

Um on the other hand yesterday we had one of our models uh discover or disprove a conjecture one of the airish problems that had smart people had worked on for a long time and a lot of people a lot of smart scientists I don't know if lun was one of them or not had even quite recently said something like that was not going to happen. Uh and then like the model just did it and you know now you have all these mathematicians saying like is math over?

What does this mean for our field? So clearly LLMs are capable of figuring out new knowledge and clearly they are capable of doing some things that some intelligence tasks that humans just can't do. Um they are going to scale much further. So how much better and what distribution of the tasks they can do better than humans. We'll find out but I suspect it's a lot. And the you know in terms of this like lack of a belief in the exponential we were talking about earlier. Um, I think the field was honestly held back by a generation of scientists who just were way too certain on what wouldn't what what scaling was not going to produce and then some people just looked at the graphs and said, "Well, it looks like it's continuing beautifully. Let's keep going." Um, I think world models are clearly important and to we'll need that for things like robotics. Uh but betting against LLM scaling at this point uh feels quite misguided to me.

>> Does it get annoying to be the I told you so guy? >> No. I mean there are these like Twitter trolls that you know for years have just been like it's not going to work. It's not going to work. This is so dumb. Like you know this is a fraud. This company's going to fail. This research approach is going to fail. And I used to get more bothered by them. But I don't even like feel the I told you so at this point. It's like you were like she was nervous.

>> You're still going on about it. Like the data is >> quite strong on our side and I don't think it'd be that fun to say I told you so. And also the fact that you're like still saying we're wrong doesn't really bother me. >> I think there's that kind of move on. >> There's that saying that like insanity is doing the same thing over and over again when presented with data that is not working and they keep repeating that. And in a sense it's it's it's a form of insanity. I think >> I I think there's something that happens which is if you make your identity about a particular thing is going to work or not work and you associate yourself with that belief and then the science or the empirical results disprove you and you're like too hung up on your identity, you can't let it go. You can't see the truth.

>> Yeah. >> And I think this is like a important reminder in both directions. >> Yeah. >> How do you see education? Um, it clearly has to super adapt and I am worried. I I thought by now it would have. Um, the the I think if we continue to teach and evaluate students as if we were in a pre-agi world, um, it's not going to work and it is going to lead to like atrophy of learning how to think or whatever. And I thought that was going to be obvious enough that I wasn't that worried. You know, when Chhatbt launched, I was like, "Yeah, we're going to have one year of like students like cheating and not learning that much. And then the educational system is just going to redesign itself." And there's and we're going to teach people so much better. You know, people are going to really get projects where they have to they have to use AI to be able to do it, but they still have to like stretch their brain more and think more and figure out new things to do. And honestly, I struggle to point to any significant systemic change that I've seen in the education system at large in the three and a half years since Chad launched.

And I that was a prediction error for me. I thought I thought that would have happened. So I have no doubt that we can uh like we have done with every other technological leap before redesign how education works so that you still have to learn how to think. And there will be some things like I I I am a person who thinks by writing and I write a lot of stuff that I never show anyone else but it's still important to me to figure something out and so I'm grateful that I I learned to write. People say the same thing about programming. Um so there will be some things that we teach people to do that machines can do better just because it's helpful to teach them the meta skill of thinking and learning and that makes sense. But there are a lot of other things where we should just totally teach totally change how we teach or how we learn or how we evaluate. And if we don't do that, I think there will be like significant atrophy in people's critical thinking skills. Uh question is what was your favorite class and what what do you wish you had taken while when you were at Stanford?

>> Does Stanford still do intro Sims? I did like all the I did like three intros a quarter my freshman year like and I loved all of them. Uh they were all super different. Uh I but looking back the fact that I was able to get such a broad exposure to stuff and h have like a a very shallow understanding of lots of different fields was an incredible thing. If it had not been for that I just would have taken like CS and physics classes which still would have been great. But um I I think more about the stuff the classes I took that were like totally random and unrelated to what I do now but in some important way gave me a perspective than I do I think I would have like learned to program no matter what. Uh so I and I didn't think that at the time I was like kind of like yeah you know this is this stuff is all cool but it's mostly going to be about like learning CS. Um, I only did two years of school. Uh, so there was a lot of stuff I wanted to take that I didn't get to. Um, but that's kind of the surprising thing.

My question is, what is your spiciest take of all? I I think with more time to think uh I could come up with a much spicier one, but um I think AI is just going to keep going. And I think this is considered I don't I don't think this is like widely believed yet. And I think if this were widely believed, there would be like significantly more reverberations that are happening through society right now. And maybe I don't have the spicier tag. Actually, maybe this is the high order bit that if AI progress continues on the exponential that it's on for another, it's been three and a half years since tragedy. If even if we're another three and a half years on that same trajectory, the world the potential the way that society what's society is capable of are just completely different. Well, let me try to prompt more thinking tokens on that one. um you you have if we treated you as a model like as a frontier model and you have some inherent capabilities and we're going we're going to try to elicit capabilities that people don't know about for the next few minutes. Um one of them is that you've been postrained now on you you've been continuously RL on OpenAI as well as the external feedback loop of the world on what doesn't work and works and doesn't work.

So now if we're going to treat you as a prediction engine for a sec, the prompt is what are the three most likely forks of the universe you see over the next 10 years and what is your what is your probability assessment on each of those? Does that make sense? One that feels very important is uh like how much is this technology going to be very widely democratized versus how much is it going to sit in a few companies. I I think a world there are all of these reasons why you could imagine the default is that this gets concentrated to a few companies and they become like you know a significant fraction of the wealth on earth that would obviously be terrible and we work super hard to push against that but I think that's going to require like the will of the world to to really avoid um because there is a sort of a tractor state there and I think part of the reason that we need to push to this kind of utility model of the world is that a it's quite unstable and quite bad will feel quite unfair if a few companies have all of this. But B, I think there's a real alignment failure and a very fragile world. Uh, and the best way to get to a world we want that represents like everybody winning and everybody's values being represented, everybody having agency is to just put push this technology out into the world.

Um, but there will be a very strong argument against that around sort of safety and stability. And I think that will be a big fork. And it's very important and I encourage all of you in your careers to push hard that this is a technology. It can bring us an incredible sci-fi future. Life can be unbelievably much better. We are going to incur some risk to get there. But the risk of keeping this concentrated in a handful of companies even though we would be one of these companies is not something we should tolerate. So I think that would be a big fork. uh in terms of probability I think it's the world should have such an interest in it happening this way that I think it's like 80% we end up on the democratic path but there will be a very strong safety message and you know there will be a lot of power seeking people who who want to concentrate the power and one of the problems with forecasting this or that you have and we all have as humans is once you make that forecast then you of agency to affect the forecasts, right? And the forecast for >> Well, I mean, we're clear on what we're going to use our agency for. Like, this is what we believe in. We think that uh, you know, we're going to do everything we can to push it in this direction. We just we see the forces in the other direction. Maybe a related fork. Uh, there's a lot of talk about like future economic models and are we going to do universal basic income? Are we going to have everybody gets to like own a slice of every company? Like, are we going to is it capitalism with no change? Is it like fullon communism? There's like a lot of talk about this. One thing that I think is not talked about much is how specifically how we distribute compute.

>> So maybe a lot of the economy can work in a way that it's going to work. And I've actually I've become much less of a even short-term jobs doomer. I've always been optimistic we find new things to do. But this may not be dup as disrupted as I originally thought in the short term. Um but we are seeing compute shortages now. I can imagine them getting much worse and I can imagine compute being like the most important utility that people need. Uh so if the price of compute from a supply and demand perspective gets way out of whack then I think there will be a very interesting fork about what it means to equitably distribute compute. So you did two very interesting things there which you said on the economic side we might have need universal basic income.

Everybody owns a piece of shares. You know, one of the speakers in this class is um Nikolai Tangjen who runs the Norwegian sovereign wealth fund. He's awesome. He's awesome. You know, the Norwegian Sovereign Wealth Fund owns 1.5% of all publicly traded companies on the planet. They also have effectively universal basic income. You could argue there's flavors of this already today because, you know, the largest employer now in the United States is the government and you could argue like large sections of that are a way for the government to redistribute income from taxpayers. So are these solutions that actually need to be novel or just reimplemented for this era? How do you think about the novelty of those solutions where we often you know in Silicon Valley make have this tendency to be like reinvent you know old things from first principles and so should we just look to existing systems and tweak them. Um yeah, I don't think that these things require deeply new ideas. Although I will say um I am much more excited about people having some sort of ownership stake than a fixed monthly cash dividend, >> right?

>> Um and I I funded like a big universal basic income study a while ago. I've also watched what happens when people like invest in startups and I know which model I think like hits human psychology better. So what I would love to see is as leverage in the world shifts from labor to capital which I think is going to keep happening that we find a way to have something like a citizens wealth fund in the country or in the world eventually where you like you basically own a slice of capitalism right a slice of these companies. And then on the second fork there on compute bottlenecks, you said uh at some point when compute prices get out of whack between January and this year, my my current understanding is based on data we've seen that H100 prices and Blackwell prices the spreads between long-term reservations and spot is like 5x.

>> I don't know if it's that high anymore. I think it got a little better. But yeah, tell me. >> Or if you can even find H100s cuz they're pretty much all gone for this year. Does that sound right? >> No argument. there's a gigantic comput shortage. Yeah. So that that's a good example of an of a systems problem right now that's live. At least to some folks, it feels like co, you know, for the comput era, like all the toilet paper's gone.

>> Yeah. >> Why are people not freaking out about this? >> Well, I think people assume we will make big inference gains on the hardware we have. Uh I also think there is a tsunami of hardware coming >> but maybe the demand tsunami is even bigger and people I think people should be freaking out somewhat >> and and would you say it's fair like how long are we going to exist in a comput shortage at least you know based on current data you have >> I think like other you you can't talk really about like worldwide demand for electricity without talking about the price like it's there's an extremely different demand about how much energy people want to use in the world if the price comes down by a factor of 10 or goes up by a factor of 10 and I think AI is

like that too.

>> Uh the if we can make models sufficiently smart and a sufficiently low cost. I think demand is like kind of uncapped and so in some sense as long as we can continue to make progress on this there will be a shortage forever and things will be bid among above what the price we think we think the price should be even though people are getting better smarter more whatever intelligence just because you can use like if we make really great personal agents and you can have 10 of them running and working for you all the time or 100 and you know you want the hundred I think >> it's a lot of inference lot of conflict.

>> Awesome. With that, I'm going to give you the swag for the class, which is [applause] Thank you for coming. Thank you. Thank you all.