

THE BROOKINGS INSTITUTION

FALK AUDITORIUM

AI AND ECONOMIC MOBILITY: OPPORTUNITIES AND CHALLENGES

Washington, D.C.

Wednesday, June 10, 2026

This is an automated transcript that has been minimally reviewed. Please check against the recording for accuracy. If you find any significant errors of substance, please let us know at events@brookings.edu

WELCOME and MODERATOR:

SANJAY PATNAIK

Bernard L. Schwartz Chair in Economic Policy Development, Senior Fellow, and Director,
Center on Regulation and Markets, The Brookings Institution

FEATURED SPEAKER:

NEIL C. THOMPSON

Director, FutureTech project at MIT

* * * * *

PATNAIK: Okay. Good morning, everyone. Thank you so much for coming in person despite the rain, and for our large online audience welcome. My name is Sanjay Patnaik, and I lead the Center on Regulation and Markets here at Brookings. Over the past six years, we have led a lot of the work on the economics of AI, so what are the potential impacts on productivity on equality and other factors, and also on how to regulate it.

And as you know, nothing much is happening in the AI space these days, right? And so today it's a real pleasure to welcome our fireside speaker, Neil Thompson, who is the director of the Future Tech Project at MIT. And we're gonna really look at a broad range of issues related to AI and economic mobility.

As you know, there is a big potential with AI to really improve our lives and have significant positive impact on things like innovation, on productivity, on new forms of economic participation. But at the same time, AI also poses quite a lot of risks. Risks that it could exacerbate inequality, that certain groups would be shut out of the progress that AI can provide.

And when we look at past technological revolutions, we have seen that alongside the innovation and the new opportunities, oftentimes it comes with disruption, with big disruption in the labor markets and with economic upheaval. And so the question is how can we channel the development of AI through policies in a way that ensures that prosperity is shared broadly and as broadly as possible, and that we can actually use it to reduce inequality and ha- make people have more opportunities than they would have otherwise.

And so we will have audience questions. My staff will be going around here with flashcards, so I please ask you to wave him down if, if you have a question and write it down and then I'll I will get them up there and can weave them into the fireside. And without further ado, I would like to ask Neil to please come upstairs, up on the stage. Thank you.

Good to have you. Thank you so much for coming.

THOMPSON: Pleasure to be here.

PATNAIK: So you lead the Future Tech Initiative at MIT. Can you talk a little bit about what you guys do?

THOMPSON: Yeah, absolutely. So, we are an unusual group at MIT. We're about 110 researchers, and we're unusual because we are about half economists and half computer scientists and engineers.

And that mix comes from sort of something that I noticed when I was doing my PhD, which was often when you're studying these kind of technologies, the actual details of the technology matter a lot for how it's going to evolve, right? And so, I sort of didn't want to do economics work where we were sort of saying, "Oh, like, let me tell you about AI today," and not be able to look into the future and think about what, what that's gonna mean later.

And conversely, on the other side, right, there's lots of work that's done in economics where it's just like, "Oh, these systems are amazing." And you say, and then people want to then extrapolate and say, "Well, therefore they're gonna be everywhere and disrupt everything."

And you say, well, actually, economics has a lot to tell us about how technologies evolve, how they get adopted.

And so what we try and do in our work is to understand what's driving AI, like, and actually where, on the technical side, what effect is that gonna have on the labor market, on the size of the economy, on science, those kind of big questions.

PATNAIK: That's awesome. And you're right at the cusp of the big wave that we are seeing, right?

THOMPSON: Absolutely.

PATNAIK: So that's very exciting. And so, part of your research is really trying to separate the AI hype that we often see—I mean, if you go to tech conferences, it looks like we're all gonna be able to live without having to work in the next five years, right?—from actually rigorous economic analysis.

And so, as you look at the current moment, what does the evidence actually tell us where AI sits in the technology adoption curve, and what does it mean for economic mobility?

THOMPSON: Yeah. So I, I think this is such an important question and you know, I, I think one way to understand this... So I think people are right when they look at it and say, "Wow, AI is evolving incredibly quickly," right? We see lots of progress, what these models are capable of doing.

But I think one thing we sort of get wrong when we think about where it is in the trajectory is we wanna infer from our experience of like, well, if something is pretty good, then it might be really, really good about it.

So let me, so traditionally, right, we all sort of know this experience of spam, right? And you were, you looked at it and you said like, "Ugh, this is, this is terrible," right? It's someone saying like, "I'm a Nigerian prince," and that you immediately saw spelling mistakes, all these things you immediately recognized what it was. And so you were like, "Oh, that's bad quality." And then there was good quality stuff, and so you had this very stark distinction between them, right? AI quickly got passable at things like that. And so we wanted to sort of say, "Oh, these seem like it's pretty good." Pretty good is not the same as perfect.

So, you know, even today in GPT-5, for example, when people, they've done the technical testing, about 5% of the time the answers have mistakes in them. Right? People have done work where they say, "Oh, well, let's use generative AI to summarize documents," and you find that sort of 3 to 7% of the time they're just adding new information that was not in the underlying document in the summary, right?

And so it's easy to sort of look at a summary and say, "Yeah, that looks about right." That is not the same as it is exactly right. And that's gonna matter a lot as we start to think about how these things get in the economy.

PATNAIK: And-

THOMPSON: The other... Oh, sorry.

PATNAIK: Go ahead, please.

THOMPSON: No, no. So the other thing I would say here about this and that I think is really important is when we think about progress, we have an intuition that comes from, from looking around in the world.

So let me, let me sort of do this by analogy. So imagine we're at the, the dawn of cars, right? We have a Model T. And we say, "Okay, it's, it's pretty good." And then we said, oh, like a year later we had a car and it was, you know, quite a bit better than that. And a year later we had a car better than that and better than that. We'd say that's amazing what the progress that is happening in cars. But sort of our intuition is like, oh, it would be kind of the same cost of car. We would be much less impressed if we said, you know, there was a Model T and the next year there was like a Lexus that was three times as expensive but better. And then the year after that there was a really expensive BMW and then a Ferrari and stuff like that. We'd say like, "Oh, there's still progress, but it's coming at a lot of additional cost."

That's really important to think about for AI because as we get these bigger models, about 80% of the improvement that we see in them just comes from throwing more computing power at them, right?

And that means that we are in that like Model T going to the Ferrari world, and it's absolutely impressive what they're doing, but we can't always just keep scaling up the cost of these things in order to do more with them.

PATNAIK: That's actually really interesting. And also your first point, I think on like not being perfect and still having the, the error rates, those can be significant for documents or for work where we need to be exact, right? I mean, and so you don't know when you get that document which information was inserted or not. The whole point of using AI is that you don't have to go through it manually. But then what do you do? I mean, if there are errors in there.

THOMPSON: Well, and, and the thing is, you know, it's even hard sometimes to know whether the system is good at this kind of thing.

So, this is an example that comes from GPT-4. I don't know if people have replicated it already in GPT-5, but if you actually ask it to do multiplication about like a three-digit or a four-digit number, it's like it's fine. It's got it, no problem. It's seen lots of that. You ask it to do it with like a 10-digit number, it almost never gets it right.

PATNAIK: Wow.

THOMPSON: Right? And so what's going on here is that we intuit, like well if you understand how to do multiplication, you understand how to do it with any number of digits, right? Like it's the same algorithm, you just keep doing it. These systems don't learn that way, they learn from examples. And so if you don't have that, we can sort of intuit and be like, "Ah, I've seen a bunch of multiplication things. It must know how to do multiplication." That's not always true.

PATNAIK: That's really interesting. Another aspect of your work is really kind of like how AI changes the economics of production innovation, right? And so when you look at today's

generative AI wave, what most excites you about the potential to really improve economic mobility and opportunity?

THOMPSON: Yeah. So, there are actually a lot of things. I, I think there's one particular thing that I would maybe emphasize here, which is that, you know, if you look at the history of economics, innovation is just incredibly important for prosperity, right? It sort of makes the pie bigger. We can then have a whole bunch of important policy debates about how to, how to share that pie, but it's really very, very important for it.

If you look at computing in general, historically, when you adopt computing, like you get a, an improvement, but actually leveraging more computing can be pretty hard, right? If I say like, "Okay, you have one copy of Excel, one computer, one copy of Excel, you can do something. Like would you like another copy of Excel and another computer to use at the same time?" And you're kind of like, "Eh." Right?

I mean, there is a small set of things that we put on supercomputers and, and things like that that make a big difference, but it's hard to get a lot of marginal productivity out of IT, traditionally. That's really not true with AI. We see that AI, because it can keep learning from more data, 'cause you can keep building larger models that have more capabilities on them, you do get this significant increase in the marginal benefit of actually using these systems.

And in, in practice for, you know, practical people, that's really, really valuable because it means that for lots and lots of problems where we said we were sort of gated by what we could do, you can imagine being able to actually invest more or get better outcomes from our systems in a way that could actually improve innovation pretty broadly.

PATNAIK: That's really interesting. And, and I think one thing that ties really closely to that is how much will it be adopted, right? It's weird. Oftentimes we only think about the technical system behind it, but yeah, it can do all these fantastic things, but it needs to be adopted, it needs to diffuse through the economy.

So when you look at adoption of AI systems, and the debate, I mean, has turned pretty negative in some quarters these days. Are you optimistic? Are you pessimistic? Do you think it's gonna be adopted and diffuse through the economy at the fast pace that oftentimes the tech companies are talking about?

THOMPSON: So I, I definitely would not say as fast as the tech companies are. I think there's an assumption there that it's gonna happen incredibly fast, and I think there's a lot of economics that will get in the way. But I do, I think it's, it's funny because I think, you know, a year or so ago, I was more, more pessimistic than most people are. And now I'm a little more optimistic than, than people are with the current down wave.

And so let me, let me explain that, right? So I think you know, there was this thought of like, "Oh, we're gonna, we're gonna adopt these. These systems are going to be great." And these predictions about like big productivity gains. And actually, if you look at the empirical evidence, it's pretty weak about that, the size of the productivity gains that people are getting. But it's not because it's not changing things. It's just that there's, there's like positive and negative happening at the same time.

So you're, you're saying like, oh, you know, if you look at coders, right? The actual act of doing the coding itself often gets significantly faster. But, maybe code review gets harder for

someone else. Or, you know, actually making sure that doing some of the review of the AI, making sure that it follows AI ethics and responsibility issues. Like, all of those add overhead to it. And so those two things have been roughly equal, and so I think that's why we're sort of not seeing much net benefit.

But of course, what we're gonna do over time is we're gonna say, "Well, how do we make sure that we keep getting more of the benefits and we develop systems that help us reduce the downside?" And so I think we will actually see more productivity gains as it comes. It's just gonna take us a little while to sort of iron out these new systems that are not really optimized because this is the first time we're applying them in many different areas.

PATNAIK: So it's a complete redesign that probably is required for our work processes, thinking processes, things like that, right?

THOMPSON: Exactly. And I, I think there's some really beautiful evidence in this that Kristina McElheran and some of her colleagues have done where they've looked in manufacturing. And what you can actually see is that as firms adopt AI, right, what you first see is that they actually become less efficient. There's this J curve where they start to get less efficient because they had incredibly optimized systems beforehand, right? And now you're saying, like, "Let's change the system." Now it's not really optimized anymore. So their, their inventories go up, their productivity goes down, and then eventually at the end, it comes, it comes back up.

And we actually see this, we've done some recent work looking at adoption in S&P 500 firms, and we see this same pattern coming where if you, for non-technology firms initially starting to adopt AI, there's some investment. You see, like, if anything, your profit margins go down. And then actually as you get sort of deeply embedded and you've figured things out, it gets, they turn positive again.

PATNAIK: And we actually got an audience question related to this, which is as we see pockets of people being quite resistant to even using AI, how do you see this play out? Do you think, like, there will be, like ways that we can bridge that, that we can bring people in and say, "Look, this is something that can help you," even if people don't wanna use it?

THOMPSON: Yeah. I mean, I, I think it's one of those things that, of course, where, where there is resistance, like, you can imagine progress being slower, but I think the gains from AI are going to be so great that, like, in other places we will see progress and then at some point people will say, "Oh, well, actually, you know, here's a way that we can do this either institutionally or technically that make things more palatable."

PATNAIK: I think that's a great point, and it really shows that we need that framework around it, right? That institutional and policy framework to make sure that we can actually deploy these technologies to the, to the best benefit.

THOMPSON: Absolutely. And I, I mean, one of the things that I imagine will be sort of a big area over the next, say, five years is trying to think about standards for AI. So that you, you can imagine many, many businesses as they're implementing it today, they sort of get, like, "Okay, well, here's the new version," right? "Wanna use it?" And you say, "Well, I don't have a lot of guarantees." You know, the, the model builders themselves have a lot of bargaining power, so the contracts don't often have a lot of transfer of risks to them. It's often sort of on, on the users to do it.

I think what we'll increasingly see are standards emerging so that people feel like, oh, they have more, you know, more security of what's going on. The insurance companies that are insuring these, these businesses for doing things will feel like, "Okay, well, we know at least people are have, you know, a model that is meeting the best standards that we can do in these things." I think that will matter a lot.

PATNAIK: And what about kind of like the downsides? So we, we have seen in the past whenever you have big productivity gains on new technology, oftentimes we have painful labor market disruption that comes with it. How do you see that play out in the AI space? And obviously there's been a lot of it written, people are trying to figure it out, but I don't think there's really kind of like a clear path forward. So I'm just curious what your thoughts are on this.

THOMPSON: Yeah. So I, I do think that there is a real potential for disruption here. I mean, I think we should, we should be prepared for that. I do think, however, that it's, it's important to separate out this question of like, is it something really systematic or is it going to be more sort of occupation by occupation?

Let me sort of say a little bit more about that. So, you know, some people have talked about like white collar work being wiped out, right? Or maybe like, you know, a lot of the expert work that we, that we do on things is gonna get wiped out, right? I think there, that, I think that's unlikely to be true. I think it's, when we look at the, the effects that we're expecting from AI, it's pretty even across the income distribution. There's, there's some changes if you do more manual work that's less addressable by AI, but it's, it's relatively flat. So it's not that it's immediately having these huge effects, but within any particular thing, there can definitely be some occupations that are gonna benefit and some that are gonna be hurt, right? And that we're gonna really have to manage that process.

PATNAIK: And, and, and what about the relationship between capital and labor? I think what, what I'm really worried about is that you will have much more concentration of wealth to accrue to those firms and, and those people that own the capital, that own the AI systems away from those people that might be disrupted on the labor side. So how do you think about that tension and what are some policies that could like counteract that?

THOMPSON: Yeah, so I think, I think this is a, a very important one, and particularly in the long run, I worry about this, right? So I think, I think in the short run, like we're likely to have - AI is likely to, to actually help the wages of some people and will hurt the wages of some people, and so I, I don't think there's sort of intrinsically that within the labor market there's a lot of effect. But of course, in the longer term, you can imagine more and more of the economy accruing to the owners of the capital because more and more of the work is being done by those systems.

And I think that's gonna be a real challenge to the way we think of equality, right? I mean, so much of like equality that we get is sort of enforced by the fact that we all have one person's worth of labor, right? And, you know, we, we get paid different amounts for that, but ultimately there's some equality that is enforced by that.

If, you know, whole parts of work disappear because we now have this move to capital, of course whoever owns that capital is going to have big returns from it. And you know, that, like managing that process. And so you asked this sort of question about, well, can we, what can we do about it? I think here, you know, things like having sort of, maybe sort of -- I think I'm, I'm sort of more open to the idea of, like, guaranteed wealth than guaranteed income.

And here just this is more about because I, I'm worried that if we have some system that says, "Well, we're gonna send checks to people every month," but it's not backed by some sort of equity or wealth stake, I think it's pretty likely that the people who have more power also end up with more political power and they can undermine that system. Whereas if you had something that is more like a sovereign wealth fund or something like that where there's actually ownership in these things, that might give more ability. You know, even better if it could be distributed.

PATNAIK: And how do you see the future of work then? Because a lot of people tie their dignity, their kind of like self-worth to work, to being able to work, right? It gives them meaning. And so what if that falls away, even if you have the baseline wealth? That's something I always think a lot about. And how do you solve that puzzle if there are not enough jobs maybe in certain areas or people don't have the right skills?

THOMPSON: Yeah, so I, I think that at least in the, in the near term, I guess I'm not that worried that we're going to lose the idea of meaning. I think there's a long history of humans being able to find meaning in a lot of things and find them very interesting, right?

And so, you know, if, if you said, "Okay, well, like, do all of the people who retire, you know, immediately feel like, Oh, I, they have no thing." It's like, well, some of them do, some of them go through a crisis, and I think, you know, that's a very realistic thing we, we should think about. But also, lots of them find out that actually there are lots of things they love that they never got to experience in their working life. And so I think we'll, we will find ways to do some of those things.

PATNAIK: And what about expertise and, and skills? So that's an audience question that came in before, which I think is really interesting. I mean, oftentimes when we look at the typical path now, how people can accumulate wealth, how can, how they can climb up the ladder, is to go to college, they go to school, they get education, they get the skills and expertise, and they get a premium for it in the labor market. That might actually go away for certain professions, right? Because certain skills, certain expertise that, that we've taken for granted, like writing, editing, things like that, are much easier with AI to bestow on people that maybe don't have it. They didn't invest four years into that. So how, how do you think about that? What, what, what does it mean to gain expertise, to, to build capacities that maybe might be meaningless with AI?

THOMPSON: Yeah. So I, I think an important thing here is to say that, you know, just because certain existing forms of expertise go away, that doesn't mean that new forms of expertise won't be created. Right? And so I think that you know, Hal Varian, the economist, has a quotation about, you know, when you have this technology that's improving incredibly rapidly, right, you don't wanna be directly in the bullseye, right? 'Cause then you're gonna get hit. But if you can be just outside of it, right, and you can sort of grab onto it and do something, you know, then it can be really beneficial.

And I think, you know, we, we know from the last, say, 30 years of technology and, and IT coming in, right? If you were, like, a dentist, but you were a dentist who knew how to use technology to do something, or if you were, you know, a physicist and you knew how to use computers to do the -- like, there was a big advantage that came from that. I think there will be lots of new forms of expertise that will be created through this.

PATNAIK: Oh, that's very interesting, and I'm glad to hear that. That, that's an optimistic outlook.

You, you wrote a paper recently called "Crashing Waves and Rising Tides," where you test whether AI capability gains arrive abruptly for narrow tasks or broadly and incrementally across the economy, and you actually find strong evidence for rising tides. So can you talk a little bit about this? Why does the shape of AI automation matter so much for economic mobility, and who gets hurt and who gets helped?

THOMPSON: Yeah. Okay, so, so this, we were, this is some work that we've been working on for, like, two or three years. So it was so exciting to, to get it out in the world. I'm glad you asked about it.

So let me, let me first of all paint a picture that I think helps me think about this question. So let's imagine we, you know, we are all workers and we're walking along the seashore. Right? And we're all, you know, we're on the, on the sand. We're feeling pretty good. We're not, like -- and here the water is going to be, like, getting hit with AI, right? So we're all walking along, and every once in a while, a wave comes along and just, like, totally drenches a couple of us, right? That's a world where, like, AI jumps forward in some particular area, and all of a sudden a whole bunch of jobs just disappear, right? That's a pretty stressful world if you don't know, like, that the wave is gonna hit you. You were thinking you're, you're good, and all of a sudden it disappears.

Another version of this of this idea is that we're all on the seashore, but some of us are, like, already up to our knees, some of us are, are, you know, are, are down to our ankles, and some of us are not. And what's happened is that the water just levels are just rising as it goes, right? And so you say, like, "Oh, gosh, I'm already up to my waist. I should already, I should be pretty worried about AI," right? "Oh, it's down at my ankles. I still have some time where I should, can start adjusting." But at least you have some sense of what's going on.

And so what we tested in this paper was, well, are we in a world where, like, we're having whole big occupations disappear, or is it more like the rising, the rising tide? And it's more like the second, right? So this gives us, I think from a policy point of view, it's really, really important because it means that we can look ahead. Right? We can say, "Ah, this occupation, like, right now these systems, you know, they're okay. Humans still have to make some changes, but the next year they're a little better, the next year a little better. Okay, let's plan for a transition of some of those workers to make sure that they're gonna be in a good place." Of course, businesses themselves can do this too.

PATNAIK: Mm-hmm. And, and I think oftentimes when we look at the debate that feeds into this is kind of like it's humans versus machine, right? But your research really shows it's more nuanced. It's like maybe it's augmenting, maybe it's complementary. Where do you see the biggest potential for human-AI collaboration, as you will?

THOMPSON: Yeah, so, you know, I think the, the collaboration can happen most seamlessly when you have a case where the, the AI is making something better and then the human can complement it. So let me, let me sort of unpack that a little bit.

So you can imagine a world where you know, we we have some task and the AI comes in and it can't solve the task perfectly, right? Maybe the human is, like, 99% accurate in, in solving this problem, and the system is only, like, 80% accurate. Well, if that, if there can be a good handoff between the machine and the human, that suddenly has made the human way, way more productive.

So let me give you a very concrete example of this. So, we did a case study with a insurance company, and they said, okay, when people get into car accidents, we want them to take a picture of their car. We're gonna feed it into an AI system, and that AI system is then going to spit out a number which says, "Ah, given, you know, the the kind of car it is, the damage you've had, here's a check." Like, you can just say accept and we'll send it to you. If you say no, then you can take it to a, you know, mechanic and get it checked out. But, like, let's do this easily.

And they tried to do this, and it was like they could not get the system to the level of quality that the humans had been at, right? That's kind of interesting, right? They would have needed a huge amount of data to get it to that level. And so they said, "Okay, well, what we're gonna do instead is we're gonna have the system pre-fill a form, and then that listed, and then the human comes along and says, 'Yes, that's right. That's right. Oh, no, no, no.'" Like it, that picture looks like the windshield is damaged. That's actually just glare from something else. It's actually fine. Right? And that turned out to be a much more effective system, but it improved the efficiency of that process by a factor of like eight or ten. And so that could be a big benefit for the, the collaboration between both of them.

PATNAIK: And I could imagine a lot of these kind of like collaborations are especially helpful in science and, and medicine, right? I know you focus a lot on science. Can you talk a little bit about that?

THOMPSON: Yeah, no, I, I, yes, it's a big deal in, in science to be able to have these systems do a lot better.

You know, one of the things that's been very interesting, we've done some work that hasn't, hasn't come out yet, but is, talks about when people are actually using, for example traditional systems like a supercomputing system or like traditional statistics for doing things, and you can see that whole new areas open up because things that used to require like a big supercomputer to do, AI can sort of simplify some of those things.

So weather prediction is an example of this right? And you can actually have much better predictions at a lower computational cost in that case. And so that's a case where suddenly, you know, our ability to use the tool has gotten magnified because the AI has changed the cost and the capabilities of it.

PATNAIK: I think the cost factor is really interesting because I think it will drive down costs for a lot of processes that were very cumbersome and have cost a lot in the past and make it much more accessible for people, right? Different types of tools that maybe no one else could have done before.

THOMPSON: Absolutely. And it's I think there's a sort of a nice analogy here. So in, in computer science, you know, there's this progression that you can see when people give people better computers, right? So it's like you were solving a problem one way, and suddenly you say, like, "Let me give you 100 times more computing power."

And often what they do is they go from like a heuristic that they were using, to like this much more exhaustive, careful optimization on what's going on. I think we're gonna see that much more generally. I think one of the things I'm very excited about for the economy is that we're gonna have a lot of new businesses that are gonna come in and we'd say like, "Oh, you know, you used to choose your school based on sort of like, ah, someone said this was a

school district that has an A rating, this was a B plus rating." Right? Now you say like, no, for my child, right, like, what are all the characteristics about the school that make it a good fit? And there's gonna be like a business around that that hasn't existed before. And that's just sort of one example, right?

There can be many, many more of these things where you can do much more exhaustive search because the, that exhaustive search has become so much more cost-effective. I think there'll be a lot of new, new businesses and new business models that emerge because of it.

PATNAIK: That's awesome. And I think given that we've seen like business creation slow down in the US over the last couple of years, that is clearly very good for the, for the outlook, I think.

THOMPSON: Very much so.

PATNAIK: You talked about kids in school, so maybe we can pick up on that. Oftentimes we get the question, and I, and I, I used to teach a class a couple of months ago at Georgetown, where students still ask like, "What should we learn?" Like, and, and how do you prepare students that are now in the K to 12 system, like middle school, high school? What are the core skills that they should be focusing on?

THOMPSON: Yeah. So I think this is, this is a very hard question and it's, you know, I sort of hinted earlier that, you know, we don't see... When we look at these things, we don't see just like, "Oh, this sector is incredibly hurt, this sector is incredibly gained, this..." You know, what we see is a much more nuanced thing, and so I think that means that it is, is actually much harder to give sort of that level of advice yet. You know, we are still working on papers that will hopefully give more, more advice on it.

PATNAIK: That's good.

THOMPSON: But I do think that you know, I would come back to what I was saying about, you know, Hal Varian's suggestion, right? Which is think about can you do stuff where you are sort of next to what these tools are able to do. And so, you know, maybe, like you say, well, I know that AI is making protein folding way, way better. Right? Something where understanding the shape of proteins allows me to answer other important medical questions, right? And so that, that sort of translation of, well, this area is gonna be improving really rapidly, let's think about the things that are adjacent to that where I can harness this tool that has suddenly gotten much better. That's what I'd point people to.

PATNAIK: That's really interesting. And I, and I do think that also a lot of the foundational skills will become important again, like critical thinking, right? Like communication skills. They will remain important. Whether you communicate with an AI or with a human, you need to be able to express yourself verbally in written form.

THOMPSON: Yes. And, and, you know, I think there though there's gonna be a, a transition where we're sort of used to teaching, for example, writing as is, and we might teach more like editing, right? Like the editor has to say, "Here's the thing you should be focused on," right? I'm assigning a, a writer to this task, and then it comes back and says, "Oh no, you've taken it off in this direction, actually, this is the direction that really matters more." And so some of those sort of editorial skills in writing and other things I think will be important too.

PATNAIK: That's really interesting. And I think, yeah, one, one aspect of this is, is your work with David Autor where you look at kind of like how job, what the impact on the job market is, and like how work gets reorganized afterwards? As you've just mentioned, like maybe the shift from writing to editing. So what determines then whether workers will be more valuable or more vulnerable as the, as the AI systems proliferate?

THOMPSON: Yeah. So this is a, a stream of work that we've been doing for a little while now, and one that I'm very excited by because I think we often have in our mind this idea that as AI comes in, it's sort of just uniformly bad, right?

You say like, oh, if a third of the tasks in your job get destroyed, it's like, well, if the demand went down by 30%, that means your wages are gonna go down, the number of people in that occupation are gonna go down. That sounds pretty terrible, right? It turns out if you look over the last forty years of automation, that is not the dominant effect that goes in. The dominant effect is something more like, well, depending on what task in your job gets automated, you have a, you have an effect, and these things can move in different directions.

So let me, let me be, be more concrete about this. So imagine you are a taxi driver twenty or thirty years ago, right? The most expert thing you knew how to do as a taxi driver was to navigate all the streets in a city, right? Sort of famously in London, there was this thing called The Knowledge that people spent you know, years studying just all the streets in London so that they can get around. GPS comes along and it automates the most expert part of that job. That means that, like, the people who are in it, actually their wages go down compared to the rest of the population, right? Their relative wages go down. But it also opens up the doors for lots and lots of people to be taxi drivers who weren't taxi drivers before, right? 'Cause you don't have to be as expert anymore. Now you can show up and be like, "Well, I know how to drive a car. I know how to read the GPS. I can now help you get from one place to another." So the number of taxi drivers doesn't go down, it goes up dramatically, right? And so wages went down, but the number of people in that went up.

What if you have the opposite effect? What if the thing that gets automated is sort of the least expert part of your job? Well, traditionally proofreaders were in this case, right? Go back not that, you know, several decades, they were actually like spell-checking what you were, what you were putting out, right? Now, of course, our computers take care of that for us. But proofread, so that, but that was like the least expert part of their job, right? Proofreaders are also like helping to craft arguments, making sure that like, oh, you know, this section is just in the totally wrong place. So their most expert part got left.

So net, when you automated the spell-checking, it was true that there was a lot less demand for what they were doing 'cause a lot of it was just like, "I don't wanna put out a press release that has a spelling mistake in it." But the ones that were left, their wages actually went up relative to other people.

And so in both of these cases, what you see is the, the employment effect and the wage effect are moving in opposite directions, right? That's much more consistent with just changes in supply happening. And so that tells us, right, that the, that when we see AI automation, we should expect in different places wages to go up, but maybe employment goes down or vice versa.

And so we, that is gonna provide different kinds of opportunities, right? When the wages go up, we say, "Ah, for some of the people in that occupation, that's pretty good. Their wages

went up." This is like a better tool that we can use. For some other cases, people are gonna be hurt who are in the occupation, but it's gonna be an opportunity for people who had less expertise to move into that. So you can imagine maybe like some nurses being able to do stuff that we only allow doctors to do right now. That could be actually pretty great for the world.

PATNAIK: Yeah. And is there a way to predict which groups or communities or workers might be affected or might be in one or the other category? That is one of the audience questions we got earlier. We already see kind of like a digital divide among different racial groups, for instance, and a rate of adoption of different technologies. Do you see that AI could exacerbate this or could actually alleviate some of those differences?

THOMPSON: Yeah, so that's, so we haven't looked in, in, in our work at, at sort of any of the differences by, by race and we're starting to look a little bit in terms of gender. I think the, the bigger thing that you-- that we can say is that it does not, it does not systematically like you could imagine, like, oh, it's just all blue-collar workers, it's all white-collar workers. We don't see that. Even within particular occupational groups, we see sort of like some people benefit, some people are hurt.

So it's not-- it isn't sort of a uniform, like pushing in just one direction or the other. It's gonna be a little more complex. That's, you know, harder to, to sort of give a pithy answer about, but maybe a little bit better for society because it means that, you know, we're, we're less likely to see the just sort of giant shifts of this gets hollowed out or that gets hollowed out.

PATNAIK: So you might have kinda like a, a blurring of the lines maybe in terms of how different communities get affected.

THOMPSON: Yeah. And it's, you know, I mean, I'm, I'm sure there will be some distributional effects that, that will come in this. I guess I'm particularly worried about that question more internationally, though. And this is just because if you imagine lots of economies where people are -- have built a, a whole industry around pretty narrow things, right? So it's like we're gonna answer customer service questions because we have English language skills, and we have this, or, you know, we're gonna do—

Like AI is going to automate a bunch of things, right, for a big economy like the United States, that means there are gonna be winners and losers, but hopefully we can build systems within the economy that will redistribute and people and figure out how they can do other jobs. That's much, much harder if you have a giant concentration of people all doing the same thing. If AI hits that, there are gonna be really profound disruptions, and that's one I worry about.

PATNAIK: So you really worry about kinda like the exacerbating inequality in between countries. That certain countries and economies that have been able to, like, lift people out of poverty and actually might lose whole industries. Is that right?

THOMPSON: Yeah, so, like at the point where we get to the country, I worry about that a little bit, but I worry about it even more at the level of just, you know, like particular geographies or particular places where you say, like, "Ah, here is a city that's dependent on this particular service that they're offering." If that disappears, the people in that city are gonna get hit really disproportionately, and you, you know, you can imagine sort of the

hollowing out of that has happened in some US cities as industry has moved on. You could imagine that happening pretty quickly in some of these places.

PATNAIK: So that's probably also bring some political disruption, I would imagine.

THOMPSON: Absolutely.

PATNAIK: And then a, a question that, that relates to that is kind of like, it's really interesting when you look at some surveys, kinda like what Americans think about AI, how positive or negative they are, Europeans and Asians, and it looks like the Americans are the most pessimistic. What are your thoughts on this? I'm, I'm very curious. I mean, the Europeans are very optimistic. The Asians are even more optimistic about it and about the ability to deploy in their lives. So how do you see that going? And, like, what do you think the reason is for that?

THOMPSON: Hmm, So I'm not sure. Let, let me ask you. I mean, if you've read the survey, what, what, like, does the survey say more about what they're optimistic or pessimistic about?

PATNAIK: It's kinda like more, I think, the, the impact on their lives, on their jobs, on, on, like, how positive the influence would be of AI.

THOMPSON: Huh. I mean, I, you know, I, I might imagine that this would, this could relate to just sort of exposure, right? Where it's like, you know, the US is sort of a, a leader in some of these things. And so, you know, maybe in the abstract it sounds wonderful and in the concrete. But, yeah, I'm not sure.

PATNAIK: That's interesting. I wanna turn to a couple of audience questions that we had before. Kinda like how -- one is, how do you think about state legislators and regulators? How should they kinda like think about regulating this? I think this is a big tension here in DC. As you know, the federal government doesn't want states to regulate too much at the at the moment, but then we have states already moving ahead with a lot of different approaches. So how do you see that tension play out?

THOMPSON: Yeah. So I guess I, I probably have less of a particular view on the state versus the federal as much as it is, I just think that there are gonna be some areas that are gonna be important to be sort of prepared for in, in a regulatory environment, right?

So we, we've already talked a little bit about over the longer term, the mix of the, the returns that go to capital and go to labor. I think that that is certainly something we need to think about. In the more short term, I think the risks of AI are a really important thing to think about. So we recently did this survey of nearly 300 AI experts, and we asked them about the risks that are coming with AI across a huge, huge number of potential things.

And not only did they think there were substantial risks, but they thought that there were significant chances of a million people dying, \$100 billion in damages, right? And this is over the next five years. And so that level of risk, I think, calls for people, you know, creating standards, creating regulations to make sure that some of these, the worst outcomes that can come from this are gonna be mitigated as they go.

I think one other thing that will be, will be important is from a sort of competition policy point of view, there are going to be some industries, and not all industries, but some industries where data and compute all sort of mutually reinforce themselves in a way that creates a, you know, real monopoly here. Right? So you can imagine a scenario where, for example, you know, I build the best AI model for something because I, I started out with the data advantage. Right? That means I have the best model, and so everyone comes to me for that.

So maybe I'm, you know, the radiology provider, right? I'll read your X-rays. And you say, "Well, why would you not send it to Neil? 'Cause Neil already has the best model." And so now all of a sudden, I also get a whole bunch more data than the next person. And so you get this sort of flywheel of I get, you know, more customers means more data, more data means a better model, means I can even more outperform my customers, right? And so in some of those places, that's going to give really, really large advantages to a particular player in a way that, you know, is gonna create antitrust challenges, I think.

Now, I do wanna say, like we should not assume this is everywhere, right? Of course. You can imagine, like in supermarkets, right? Like sure, more information is useful. It's still the case that the corner store that is two minutes from your door rather than the store that is an hour from you, like they don't necessarily need all that information to be able to serve you pretty well and for you to still to keep going there.

And so there are lots of other things that matter, and we see, you know, historically, very large differences in productivity can still persist in industries. So it's not that it's sort of gonna immediately, but in some places it could be a real challenge.

PATNAIK: And so how would you approach it from an antitrust perspective, given that I think that would change a lot, I mean, a lot of these industries integrate across sectors these days, so it's already a challenge for regular antitrust enforcement. When we look at current sectors, like, Amazon is active in so many different industries. How do you think that could change with AI, even the antitrust approach for regulators?

THOMPSON: So I'm not sure I feel like I know enough about the antitrust toolkit to, to, to know the best solutions for these things. I mean, you know, like certainly you can imagine a scenario where you'd say, well, like, how, how can you make sure that other people have good models in the – and so, take a, I mean, if you think a example inspired by telecommunications, right? You could say like, "Well, we allow multiple people to use this model," or, "We allow, you know, we wanna make sure that, you know, you have to create two companies and each of those companies has access to the model," or something like that.

But, you know, like these are things that would carry with them significant disruption, and so I don't wanna just sort of like advocate for them. I think you can imagine options here. But actually, what is the best one I feel very unsure about.

PATNAIK: Another question that came in is, I think that's an interesting one. A lot of critics think that this actually might be a repeat of the dot-com bubble, all the money that is pouring into like all these data centers and all, as you mentioned, into the compute. What's your view on this? Is, is this a bubble or is it not? Hard to say, you probably would be an oracle if you, if you knew that.

THOMPSON: So I guess, let me, let me split this up into like is it a financial bubble and is it sort of a, like, is it just all hype? 'Cause I think those are sort of separate issues.

So I think, I think it's very hard to know the question about whether it's a financial bubble. And the reason for that is just there are so many things that are moving so fast in different directions that what some of them are gonna overwhelm others and exactly which ones are unclear.

So if you think about things pushing in the direction of like, actually, yes, we need these data centers and they're big. It's, you know, adoption is going up very fast, and the size of models we're using and the amount of computation that they're using is going up very fast. So all of those things are pushing for lots more demand and, and on the other hand, we also see chips are improving pretty quickly. The algorithms that we're using are improving pretty quickly, and all those things mean that over time, like you get to a model and you say, "Ah, you know, now I'm, I'm a doctor's office. I use it to transcribe." Maybe a year from now you can use a model that is one tenth as big to do that same thing. That's actually gonna save a lot of compute, right?

And so exactly how these things play out, I, I think that's a actually a very hard question to answer. However, I, I feel quite confident in saying, like, this is not -- like, at the bigger question of is this hype? This is not all hype, right? There are very, like, even in cases where we might see, like today you might -- so what some people call, like, so-so automation, this idea of like, okay, the AI takes the job that the task that the human was doing, but they're sort of comparably efficient in something. And you say like, "Well, we lost a full job and we didn't get mu-much of a gain from the AI system." Like, that's not, we sort of lost a lot and we didn't gain a lot in return.

The thing is, because these AI models are improving so fast, also in terms of efficiency, in terms of hardware and stuff, anything that's so-so today will not be so-so later. And that means that over the longer term, I think we should expect lots of improvements, lots of productivity benefits to flow. We just have to go through that process of figuring out how to make these things work, how to capture the benefits with the most efficient model, all of those things. We will mature to that, I think.

PATNAIK: And how about the, the way firms deploy those models? So one question here, it talks about the risk when, when you have increasingly autonomous systems, which we see everywhere, how can firms manage risks associated with that? I think it's gonna be tricky, right?

THOMPSON: Yeah, it, it is absolutely gonna be tricky and I was on a, a panel with someone from one of the big AI producers, and, you know, one of the things they talked about is which I thought was very interesting, was that, in some cases where there was a, a particularly lot of risk, they would actually run, basically companies would run two systems in parallel. So they'd have like the AI system, which is usually better, but you wanted to make sure it didn't go off the guardrails. And you had a traditional system which wasn't as good, but it, if the AI system ever diverged too greatly from the other one, they would default back to the, the traditional one. And so that you had this sort of guardrail of like you can't go too, too wrong with it. So I think that's sort of one of the things you can do.

But in terms of the more general thing, I actually think it depends a lot on the specific risk. There are a lot of different AI risks and a lot of stuff. And so I, I would point people to my lab has this database. You just go to airisks.mit.edu. And we have right now a big taxonomy of

all the risks. In the next year, we will have also a taxonomy of the mitigations. And so for each particular risk, you can actually think about what is the thing that you, you should be doing and what's the most effective thing you should be doing to mitigate the risk from that.

PATNAIK: That's awesome. One thing related to this is, and we've been thinking a lot about this, and this is another question that came up, is you have to disconnect with a market and technology that moves incredibly rapidly, as you mentioned. Like in three weeks already the model looks very different. But we have policymakers and regulators that are very slow.

And they have been slow over the last couple of decades, and now it exacerbates that disconnect because by the time they come up with policies or regulations, they're already outdated. How do we bridge that gap? How can we like make sure they can respond actually in almost real time to what we see in the marketplace?

THOMPSON: So I, I wanna argue that for a lot of this stuff we shouldn't try and be real time. We should try and understand the deeper trends so that we can see what's coming and, and regulate for that. Right?

And so, you know, I sort of started this discussion when you asked me about what FutureTech is, you know, I, I talked about not wanting to write an economics article where you say, you know, "Thank you for telling me about what AI looked like two years ago. That's not what it looks like today." I think we can say the same thing about this, right? What we want is to be able to create laws and regulations that are not for, you know, somebody's anecdote about what AI was like two years ago. We wanna know, like, okay, here's the trajectory that these systems are on, and therefore, you know, this is gonna change, this is gonna change. Let's prepare for those things today. I think that is possible, but you know, it's, it's sort of complicated because you have to understand some of the technical pieces going on.

And so, you know, I hope that more and more people are trying to think about that trajectory, not just the immediate moment of where things are with AI today.

PATNAIK: Kinda like find a baseline level of regulation that can accommodate multiple different pathways.

THOMPSON: Yeah. Well, and, and so not even necessarily multiple pathways. So let me, let me give you a concrete example of this, right? I mentioned this idea that 80% of the improvement that is coming from these AI models, the frontier models, in the last little while, is coming from just throwing more computing power at it. Right? We know that if you just sort of extrapolate that forward, within three to seven years, right, you get to very unsustainable amounts of money, right, for it to keep scaling up these models, right? So that means that these systems are gonna change, and you can already start thinking about that ahead of time and saying, "Okay, well, you know, I, I can think about m- what the model's gonna look like when, you know... I can think about that that's gonna put pressure on people to take models and try and distill them into smaller models."

Like, you can sort of predict some of the things that are gonna happen. Now, of course, you know, this is across a broad range of regulations, so you'd want to think about it for each particular case. But you actually can make some projections about what these systems are gonna be capable of and therefore, and how they're gonna evolve, and that means that you can regulate them in a sort of more intelligent, more capable way.

PATNAIK: That's actually very helpful. I wanna finish with one question, which is interesting. A lot of people have compared kinda like AI and the race that the US finds itself in with China and other countries as a core national security issue and, and compared to the past almost, to the Manhattan Project, right?

And so there are ideas out there that, that propose that the US government should take a stake in the AI companies and help with the development or channel the development in certain directions that can strengthen US national security. What is your thought on that? That's a very different approach, but we have seen it in the past, like things like nuclear development and others that are relevant for our national security.

THOMPSON: Yeah, so I think this is it, it's an important question, and I wanna sort of give a, a little bit of a sort of middle answer on it. So one of the questions that my lab has asked is, you know, does building the-- a larger model, like, give you more of an insurmountable lead or less of an insurmountable lead?

Now, the, the why -- so why do I bring this up in the answer to this question? Well, you know, the thing you would most want to do, the reason you would want, most want to do this is, is to say, like, AI, these AI systems, like, if we think we're in a competition, right, we wanna, we want to sort of our system to be better. Right? Because we feel like it's going to give us a, a distinct national security advantage or a distinct advantage in something else.

And then the, then the question is, like, okay, well, that may be true, but to what extent, right? Like, is this something where if we don't do this, there's gonna be a one percent difference between our systems and Chinese systems in terms of what they're able to do? Okay, we might not worry about that 'cause we think that one percent is gonna be swamped by all of the other stuff that's going on in national security. Or is it gonna be, like, a fifty percent difference where we'd say, like, there we really care about it.

And so, in some work we did where we, we talked about this in terms of meek models, how strong are, are these meek models? So if you say, like, "I'm gonna train a model and it's only gonna have a thousand dollars. I only got a thousand dollars to build my model," how different is that than the model that you can build at the frontier? And what you see is that in lots of areas that you actually work in of course, the big model gets there first, but the meek models catch up over time. And actually the gap between them actually narrows over time.

So if you think about, like, you know, performance on some benchmark and you say, "Okay, I want, you know, like, I wanna get to ninety percent on this benchmark," the, the big models will get there first, and then the gap between them will shrink and shrink and shrink as the, the, these small models become more capable.

And this is you know, there's some deep technical stuff going on here about why this is true. So the, the places where this isn't true is where you get really unbounded kind of competition between stuff. So if you think about, like, okay, well, you know, strategy for this military versus that military, then you actually can get situations where the bigger model sort of runs away and gets, continues to get more performance and stuff and thing. Okay. So I think there are, there's at least enough reason here to think that, like, this is a very serious worry. That, that it's—we should absolutely be thinking seriously about this question about how to promote, like, like, make sure the national security implications are good. But you know, whether this is exactly the right mechanism, I feel unsure about.

PATNAIK: Great. Well, this has been fantastic. Thank you so much for being here today and for your great work. Thank you so much.

THOMPSON: My pleasure. Thank you.