

*****Please note that while these transcripts were produced by a professional, they may not be entirely precise. We encourage you to use them for reference but consult the video to ensure accuracy.*****

ABCDE 2025

Tuesday, July 22nd 2025
Washington, DC

ARTIFICIAL INTELLIGENCE

Karolina Oordon: We're about to start the second session today, which is all about the future, specifically how artificial intelligence is revolutionizing technology, reshaping competition policy, and redefining education. We'll hear from experts about the opportunities and risks of AI for developing countries, and how policymakers can use this technology for inclusive growth. Moderating the session is Mr. Gaurav Nayyar, Director of the World Development Report, 2026. He'll introduce our three panelists and one discussion, so over to you, please.

Gaurav Nayyar: Hi, good afternoon, everyone, to the post-lunch siesta session. But I'm hoping that just the mention of the word "artificial intelligence" will wake everyone up. So thank you for being here. The esteemed guests, panelists had a preference to not be introduced as rock stars and enter the stage one by one with a spotlight light on them. They've all taken their seats. I just want to quickly give a very, very quick introduction to the session. In the morning, we heard about populism, and populism as a reaction to distributional concerns both between countries and within countries. I think the discussion around AI dovetails almost perfectly into this discussion on how the diffusion of the technology might affect the distribution of income, both between and within countries. Here at the Bank, I think there is a lot of optimism around the development potential of AI, just in the sense of the technology being able to solve market failures and addressing long-standing development challenges. Think about access to credit markets where AI can devise new ways of assessing credit histories. Think about filling in skill gaps, say in education and health services where quality teachers and health service providers are often lacking, or just making the large numbers of small businesses more profitable through AI advisory services.

However, such optimism is not unbridled. I think everyone recognizes and is aware of the risks that come through growing digital divides, competition concerns, and also risks of automations and worker displacement. The aim of today's session is really to cut through some of the hype around AI, provide some conceptual clarities around specific sets of issues, and then also present some new evidence which will hopefully give us a better sense of the potential development impacts of AI going forward. To do that, we have a stellar set of panelists. To begin with, we will have Arvind Narayanan, who is Professor of Computer Science at Princeton University, talk about AI as a normal technology. After that, we will have Catherine Tucker, the Sloan Distinguished Professor of Management at MIT, talk about how does competition policy need to change in a world of artificial intelligence. Last but not least, we would have Dan Björkegren from Columbia University, who will talk about could AI leapfrog the web and provide evidence from teachers in Sierra Leone. Also, immediately after the presentation of the three papers, we will move into a panel discussion where we will also be joined by Han Sheng Chia, who's a Policy Fellow at the Center for Global Development.

With that, I'd like to invite Arvind to give us the first presentation. Thank you.

Arvind Narayanan: Thank you, Gaurav. Hello, everybody. Good afternoon. My goal for the next 20 minutes is to say things that will sound obvious to you, hopefully. That's an unusual goal for a talk, so let me explain why that's my goal. Perhaps if we could have the slides up in the meantime. Now, there are many technologists, AI researchers, developers, and industry CEOs who have predicted imminent and rapid transformation because of AI, that we will have super intelligence in the next couple of years before the end of the decade, certainly, and that this might obviate the need for human labor and perhaps even the concept of money, and that we should be worried about AI taking over the world. The more modest of these predictions are talking about 10% growth annually. But as economists, a lot of this might seem strange. It seems to be from the realm of sci-fi, so it might be natural to wonder, do these folks know something that we don't know? I'm here as a computer scientist to say, I think the answer is no. I have looked at these predictions and the research behind them. My view is that, together with my co-author, Sayash Kapoor, that AI is not going to be that different from many past waves of technological disruption that we've seen.

Of course, no two technologies are exactly like each other, but there is a lot that we can learn from history. That's what our paper, AI as Normal Technology, is all about, where expanding it now into a book. The paper has four parts. The first part is about timeline, on what time scale will AI's transformative effects unfold? The second part is about, "Okay, whatever time scale that is, what will that world look like?" Now, we accept the possibility of cognitive automation that most of the cognitive tasks that we do today might one day be automatable. And yet we do not envision something radical such as human labor being superfluous. Why is that? What is the resolution to that seeming contradiction? The third part I'm going to mostly skip for today, what do we do about potential superintelligence, etc. Then I'll briefly spend a couple of minutes on policy implications, and I look forward to the discussion. Our starting point, as well as tabulate established patterns of the diffusion of general purpose technologies that we've seen over and over. The first stage is invention. If we think about electricity, for instance, the invention of electromagnetic principles, AC and DC, those are all part of invention.

But people don't directly use electricity. We use electrical appliances. So that requires innovation. It requires the grid. It requires a whole lot of other complementary things. Then finally, that technology gradually diffuses throughout society and throughout economy. We took this and we expanded and adapted it for AI. We introduced one extra stage, and we talk about what each of these stages look like in the context of AI. The first invention stage is about the increasing capabilities of models. Today, the model is the model that's at the frontier are large language models. It wasn't always the case. It might not necessarily always be the case, but for now, you can think of the invention stage as improvements in LLM capabilities. We don't necessarily directly use LLMs for most tasks. We use it through the integration of those LLMs into various applications. So throughout this slide, I'm using software engineering itself as an example because it is one occupation that's really at the forefront of integration of AI into workflows. So even though software engineering is pretty advanced along this causal chain, if you will, we're still pretty early in figuring out how to architect AI code editors.

We're having really fundamental debates in the computer science community about whether command-line interfaces are better or graphical interfaces are better when it comes to software engineering using large language models. It seems like very early days. Then you have the early adoption phase where people are learning how to use these tools. That phase is very much underway. You have things like vibe coding where people are able to create small apps using AI. But when they start to edit billion-line enterprise code bases using AI, it's not going well so far. We're not yet in the adaptation phase. We're not yet in the phase where people are figuring out how to re-architect software, both technically as well as software companies, to best take advantage of the capabilities of AI tools. There was a study that showed that right now, most software engineers are not yet achieving productivity benefits, even though they think they are. Their self-reported time savings are actually very positive, but actual reported time savings are negative. Our view here is that adaptation is going to require a lot of changes. One example is what I'm calling extreme personalization. This is all well into the future, so it's necessarily speculative, but let me explain what I mean.

If it's going to be the case that AI's ability to write code will dramatically decrease the cost of producing software, then the manner in which we're going to use software is going to change dramatically. It will no longer make sense for one software company to produce one piece of software and then millions of other individuals and other companies to all be shoehorn into that same workflow, the same set of abstractions that's defined into that software. The way we produce software might be in a very in-house way. Each company is producing its own software for its own teams and for its own users. In other words, if this cost of software drops, the elasticity of demand for software might be very high. We might see extremely personalized software where everybody is using software tailored to themselves. If you think about that magnitude of transformation, hopefully it's intuitively clear why that feels like a few decades away as opposed to something that's going to happen next year because of increasing AI capabilities. That's really the central idea of the paper,

why we think these timelines are going to be pretty slow. Okay, so let me focus, though, on the first era, the gap between the model capabilities and the products.

Now, in the first year after ChatGPT was released, many companies tried to create these simple wrappers around these large language models and build these personal agents, but they failed terribly in the market. There's a whole graveyard of products like this. Here's one where this product does a pretty good job of automatically ordering food, so it's an impressive capability, but it ordered it to the wrong address. We've had unreliable AI like this before. If traditional Siri or Alexa played the wrong song 10% of the time, that's just annoying. But if a product orders something where you spent money to the wrong address 10% of the time, that's dead on arrival. This is something we call the capability-reliability gap. We think there's years of research to go before we smooth out these kinks and build reliable products using large language models. Okay, so capability-reliability gap is one reason why the way we're evaluating AI today doesn't make sense. We're evaluating models and assuming that that means useful products, but that's not the case. Another reason is construct validity. For instance, there are studies that AI can pass the bar exam, but it's not as if lawyer's job is to answer bar exam questions all day.

The evaluations are looking at the wrong construct. We don't have a lot of information about how useful these are to actual workers unless you do RCTs of those workers using those tools. We have lots of anecdotal evidence now that the hype doesn't match the reality. The RCTs in many domains are still just getting started. This is a company [that] do not pay that got a lot of attention by saying they'd built a robot lawyer, and any attorney who used it to argue in front of the Supreme Court would get \$1 million. First of all, this was just a publicity stunt. Electronic devices are not even allowed in the Supreme Court. Of course, this company knew this, but the reason so many people believe this was even possible was because they believed the headlines and they thought that a robot lawyer was imminent. But when you look at lawyers actually using AI to do the hard parts of being a lawyer, it's still very, very far to go. We teamed up with our colleague, Peter Henderson, who's a JD PhD, and we wrote a paper breaking down the different types of tasks that are involved in being a lawyer.

It turns out, and this is the central reason why so many technologists are overconfident, the tasks that are easier to measure, which is the top bucket, are also the ones that are much easier for AI. We just measure the things that are easy to measure and conclude that AI is doing really well at it. Okay, so far I've talked about the first era, but the much harder thing is diffusion, the second era. Let me talk about that a little bit. One important point I want to emphasize is that it doesn't matter how fast capabilities are increasing. Adoption of these technologies requires behavioral changes, and behavioral changes are going to happen at a certain time scale based on how fast humans can adapt, and that's not really changing very much. There was this paper that made the headlines. It's a good paper. I have no concerns about the paper. But the thing that the media picked up is that 40% of US adults are using generative AI as of about a year ago, so two years after ChatGPT. That seemed very striking and it's framed as rapid adoption of generative AI. I don't think this tells the whole story, because what the paper also measured, excuse me, is the intensity of adoption and productivity effects, and those are still very low.

There is still a very long way to go because That 40% number includes even someone who used ChatGPT once a week to write a limerick or something. It's not necessarily looking at how much productivity improvements people are getting in the workplace. Okay, and think about this. I mean, so often, and I'm part of this, I often forget that Google Search exists and things that I'm debating with other people can be quickly resolved through a Google search. And this person is telling a story of how much of a competitive edge you get just from remembering to Google the things that can be googled. And that's a quarter century after Google Search was released. So that's the time scale on which these behavioral adaptations happen. Okay, so that's on the speed of diffusion. However, one of the things we point out in the paper is that in the last few years, we've been in an unusual period

of AI innovation, and that has already ended. What I mean is this. The reason why we had such rapid AI capability advances was building bigger and bigger models using bigger and bigger slices off the Internet. That has mostly ended, not completely, but for the most part.

What it means going forward is that these models will have to be improved using experiential data from their use in the field by actual workers. That's going to be a feedback loop that's going to happen much slower, and we're already seeing those much slower cycles of progress kick in. We can reasonably predict what that will look like by looking at sectors in which those cycles of feedback loops have already happened, notably self-driving cars. What happened in self-driving cars is that two, three decades ago, we had convincing prototype demonstrations of self-driving cars. However, they weren't yet reliable enough to autonomously put on the roads. What car companies had to do or AI companies had to do, was to get it reliable enough to drive 100 miles, drive it 100 miles, get enough data, improve their reliability so that you're comfortable driving it 1,000 miles, and so on and so on. Each order of magnitude of scaling takes a couple of years. We've seen those slow feedback loops play out, and we predict the same thing is going to happen in AI for law, AI for medicine, and every other sector. Finally, what we say about the speed of change is that a lot of this might require changes to organizations, structural transformations, changes to even laws and norms and so forth.

We use the analogy of how slow the electrification of factories was because they had to figure out how to reorient the whole assembly line and so on. Looking at that example of software again, so if it's going to be the case that it makes no sense to have these large software companies, and every individual company is going to have its own software engineers, that is a slow transformation of the whole sector that is, again, going to play out on a time scale of decades. Okay, so there's lots of other examples of this, but really the summary of the first part is that we identify speed limits in every stage of this invention to innovation, to early adoption, to diffusion pipeline. I haven't gotten into the details of most of the speed limits, but they're in the paper. And that is what leads us to conclude that we're not going to see some sci-fi transformation in 2027, but rather we're going to see gradual diffusion throughout the next few decades. So let me spend five minutes on part two and then maybe two minutes on part four because I know I'm running out of time. So again, like I said earlier, so we accept the possibility that it might be the case that one day, most cognitive tasks that we do, most things we do sitting in front of the computer are going to be automatable.

But let's think about what that will mean. There have been concerns about mass job displacement because of automation for many years. This was a headline from 15 years ago. However, from what I understand of the literature and economics, those are naive ways of thinking about things. There is the task model that goes back to Autor, Libby, and Mournay[?] more than two decades ago and has been refined over the decades. We can see jobs as bundles of tasks, and each job might have several dozen tasks as part of the bundle. Not every task gets automated at the same moment in time. As one task gets automated, new tasks get created, the job composition changes. Automation is counterbalanced by a reinstatement effect. It could even be the case that if it becomes cheaper for lawyers to provide their services because of automation, there's more demand for legal work, and therefore, the total employment of lawyers actually increases. Okay, so that much, hopefully, is straightforward. But now, I want to make a couple of points that are even more skeptical of mass automation compared to what you might take away from the literature in labor economics. For that, I want to give you an example.

Radiology was famously predicted by computer science pioneer Jeff Hinton, a Nobel Prize winner, to be going obsolete, because AI can automate the reading of scans. He said almost a decade ago that if you work as a radiologist, you're like the coyote that's already over the edge of the cliff but hasn't yet looked down. In my world, this is a famous prediction because of how wrong it was. Demand for a radiologist has only gone up. Even though radiologists are not shy of the tech, they have enthusiastically adopted it, and it has had an effective augmentation rather than automation.

Interestingly, this does not appear to be because there are some particular tasks that we can identify at which human radiologists still outperform AI. This is the central claim that I want to make. I think the task model of jobs is incomplete. Yes, to a first approximation, you can break down jobs into tasks, but think about your own job. Can you identify a series of tasks such that if each of those tasks was automated, then your job as a whole would be automated? I would say no, and I would imagine that most of you would say no, and you might even find the question to be a little bit silly.

And yet that's exactly how I think a lot of this literature approaches the question of AI and job displacement. In my understanding of radiology and many other professions, a lot of the hardest to automate things are at the boundaries between tasks. So by the time you've broken down jobs into a set of tasks that you can put into the own that database that everybody seems to use, you've already lost the most nuanced and interesting aspects of jobs that are going to be hardest to automate. And then you have to consider human preferences for dealing with other humans instead of AI and various other things. So considering all this, we're very skeptical about some of these claims of broad-based labor displacement due to AI. That leads to the point, I think there are a lot of these predictions made by computer scientists like us, made by economists who are not necessarily experts in the specific domains that they're making predictions about. I think there is some value in those. I mean, it is a lot of what I do, and I'm very much turning this criticism back onto myself as well. But I think it would be much more useful if we teamed up with experts in different domains who have a much more nuanced understanding of what tasks entail, what are the boundaries between tasks and so forth, to make predictions and recommendations about those specific occupations and sectors, as opposed to these overarching sweeping pronouncements.

Okay, so having said that, I'm going to return to the mode of making sweeping pronouncements. One thing I want to point out is that there is actually a way to reconcile some of the radical optimists and the radical skeptics. We can do that by looking back at the effect that the Internet has had. Now, let's accept for a second that generative AI will not this is early automated, but at least mediates most knowledge work, just like the Internet did. If you went back in time 30 years and told people that most of the cognitive work that they do will one day be mediated by this new communication technology, they would have said, "That's ridiculous." A lot of people did say, "That's ridiculous." And yet that happened, even though it sounded crazy at the beginning. What were its effects? The effects on productivity, I would say, have been relatively small. And while it's true that some professions have been decimated, for the most part, we see an augmentation effect. And this all unfolded gradually over a period of a few decades. And therefore, our thesis is that you can have workplaces that look in some ways radically different, and yet without necessarily having transformative effects on total factor productivity, growth, or something like that.

I'll skip this slide. To be a little bit flippant about it, our view on the role of labor in a world with advanced AI is that even though everything might change in a qualitative sense in terms of our experience of what work looks like, in an economic sense, things might mostly stay the same. One thing I want to add to this is whether superintelligence will change the equation. A lot of people who have very different predictions from us will say, "No, this is what past technologies have done," and so far AI is on the same path. But something is going to change in a couple of years where recursive self-improvement of AI becomes possible, and then it will be an exception to all of these well-established patterns. We have a whole section in the paper on this, but I'll just spend one minute on it. We say no. The reason we say no is that we have a very different view of what human intelligence entails and to what extent AI will be able to exceed human intelligence. You can identify three types of limits to our capabilities. We suck at chess, and AI is much better than even the human world champion because of our computational limits.

But for most tasks, the limits are not computational. There are things like writing where there are intrinsic limits to how persuasive let's say, a piece of writing can get, or there are knowledge limits. Some people say AI is going to be a superhuman medical researcher, and we're like, "What? Let's

take a second to imagine what that means. For AI to be a superhuman medical researcher, it can only happen if we allow AI to perform autonomous experiments on thousands of people without any supervision." We will never allow it to do that. There are intrinsic knowledge limits to how good AI can get. We reject this notion that superintelligence is going to introduce a new causal pathway that's going to blow past all of these limits that are identified from previous research. Okay, so I will stop here in the interest of time. We have lots of thoughts on what policymakers should do, but maybe we can get into that during the discussion period. Thank you so much.

Catherine Tucker: Could I just say how brilliant that was? We got a computer scientist summarizing all labor economics and AI for the last 10 years. It was really well done. Is there a chance of my slides? Brilliant. Okay. Now for something completely different because I'm an IO Economist, which means I've been tasked not to talk so much about labor, but to talk a little bit more about how AI will change, how firms relate to consumers, and how that will change in a world of AI. Now, I was asked to talk about a paper I wrote on this topic for this panel, and I must admit, when I was asked this, I was like, "Oh, that's my paper which has got least to do with any low income countries." It's completely obsessed with America, completely inappropriate for this audience. So what am I going to do with the fact that this paper is very US-obsessed? What I'm going to do is I'm going to set out the broader lessons, because in some ways, the reason I wrote this paper was because much like Arvind, I often sit on panels with lots of policy people who have very dramatic perspectives about AI.

You sit there and the general view about competition policy and how firms relate to consumers and how everything's going to change is everything is going to change very, very swiftly. Then there's always one person, a little bit more like Arvind, who says, "No, I'm going to be different." Then I'm going to be more economist and be slightly in the middle and say, "Look, in the end, economics doesn't change." That's good news if you're an economist. On the other hand, process change is going to be important. Now, the great thing about writing this paper for me was that I wrote it in response to being on a lot of policy panels and really taking it as a chance to try and educate lawyers, largely about how economists think about artificial intelligence. I have run the NBR meetings on digital economics and AI for nearly a decade now. I want to just take this opportunity, not just given the article, to just tell this group a little bit about the history of the economics of AI. I think it's fair to say we started off back in 2018 with our first NBR official meeting about the economics of AI.

What was the revelation back in 2018? I think the revelation for economists back in 2018 is that AI was not about robots. I think we never really written many papers about it because we thought it was about robots and therefore about Terminator and Judgment Day. Instead, what we worked out was instead that really, if you think about AI, it is a cost-saving technology. The moment you start to describe AI as a cost-saving technology, it is something which economists start to think they can analyze, unlike Terminator. The fortune back in 2018 is the way to think about it is that, AI lowers the cost of prediction. As you may have guessed from Arvind's talk, who talked about Hinton, what he said about radiologists back in 2016, we were obsessed about radiologists. And where we came into the radiologist debate is we said, and we thought we were being very smart, is that, "Look, the prediction is that radiology is going away." But in the end, as economists, we believe in complementarity and the complementarity of tasks. And therefore, really, what's going to happen is when you reduce the costs of doing certain type of tasks, such as prediction, other forms of humanity are going to be more important, and the one we emphasized there was judgment, or the idea that there's going to be large returns to judgment in our AI economy.

Now, at the time, we were all joking around, I remember quite a few laughs were had at our 2018 conference by someone suggesting that what we should all be doing is encouraging our children to become journalists and poets because those occupations were not going to be disrupted. What I learned from that is that economists who study AI are usually wrong because now in 2025, that viewpoint that AI is not coming for creativity seems completely wrong, especially if you're a poet.

What I've learned to... When I talk about the economics of AI more broadly, especially to PhD students, what I've always tried to do now is to encourage them to be a little bit more humbled, express humility, don't have our hubris that we had back in 2018. And as I returned to this Oxford English dictionary definition, which is that in the end, what is AI? It's the set of systems and technologies which is capable of performing tasks which require human intelligence. And what I want you to notice is that is a moving target. That definition can never be true or stable in the long run, because the nature of AI is such that, of course, we're always going to be finding new ways of replacing human intelligence.

So I've learned to be very humble about how I talk about AI. And that was one of the things I was trying to get across in my article, which was really a debate, as I say, with antitrust lawyers. Now, the one thing that all AI economists do agree about is very much, and this is so similar, and that's why I was congratulating Arvind on his talk, in that the one thing we all agree about is that AI is probably a general purpose technology. That makes technology economists happy. Why? Because we have so many examples littered around. And as[?] you'll have gathered, one of our favorite examples is electrification and making that analogy. All I want to really add to this very good evidence that we just heard in the previous talk is to remind everyone in the room that when electricity was first introduced, everyone thought it was going to be about illumination. We all thought that the jobs that were going to be gone are the people who around the streets lighting gas lamps. We had no idea of where the real impacts would come from, which was the electrification of factories. Again, in 2025, I think we're probably at that stage in that we're still really thinking about use cases which are the analogy of illumination rather than really reflecting fundamental process change, which is what underpins general purpose technologies.

Okay, so I got to set all this out in the article about the history of the economics of AI and how economists have tended to think about it. Then I realized that what I'm really doing is I've set myself the task of writing an article about a technology which is fluid, unpredictable, has no established business model, and bring it to competition policy. If you know anything about competition policy, there's a big mismatch here. Why is that? Well, in the end, competition policy is the belief that governments do good by encouraging competition, and that competition is a way of benefiting the citizenry and the average consumer, and that's the aim of competition policy. Now, the reason I say that the unpredictability of AI poses trouble for competition policy, is that in the end, the way that competition policy tends to work is that it tends to be very retrospective. If you think about US competition policy, often it is implemented through the courts, and you are therefore always going to be looking backwards. Again, we've got this technology which is changing its nature every two years. The other thing I'd say about competition policy is across the world, there's generally a set of norms about how you conduct it and how we think and analyze substitution by people across products and firms.

There's also a lot of norms of which are established in the competition policy tends to be very obsessed with looking at price. If you look at all the tools that IO economists tend to bring to competition policy, they require you to look at a price. Now I've got to say something about technology where we literally don't know what the business model is. We have lots of tech companies spending billions and billions of dollars on it with no return, giving a lot away for free. Again, that is going to make it difficult to predict. Now, when I said I was going to give this talk, I inwardly panicked and reached out to Daniel. Why Daniel? Well, you'll hear in the next talk, Daniel's a brilliant development economist. Good news. He's going to talk about AI in developing countries. Again, good news. I was like, "How do I say, how do I bridge this more to a group of development economists?" He was like, "No, no, Catherine. You're really underestimating the extent to which antitrust is important. For example, look at these evolutions in Africa." I'm going to give credit to that and let Daniel talk about that.

But what I would just say is that I think this debate is still important about how we regulate competition around AI, even if at the moment it has a US focus. The reason is this, if we were to look backwards, as competition economists often do, at what's happening in antitrust and technology more generally, decisions being made about competition policy in the US are still going to have implications for the rest of the world. I'll just give you an example of both Google and Meta. Both of them, I know we're probably moving in completely parallel worlds in DC right now, but there've been two really large antitrust cases for both of these companies in DC in the past year. I would argue that those cases are going to have implications worldwide. Why is that? Well, if you think about current tech platforms, there is this I think less-talked-about aspect of them is that they tend to be worldwide, and there's an awful lot of cross-subsidisation involved. Basically, they are built around business models that thrive away from the fact that rich people's eyeballs are very valuable. If you put ads in front of those rich people's eyeballs, you can make a lot of money.

The products, though, are worldwide, even if all the have been generated in a few rich cities. Therefore, when antitrust policy or competition policy comes in and tries to alter those business models, there might be some more worldwide implications. Okay, so I've got five minutes, and I should say a little bit more about my views about competition policy and how it should change in the world of AI, given that's what the title of the talk was, or the paper was. Again, my message, as I said, to summarize, is that economics doesn't change. We're still going to have exactly the same digital economics concerns that we've had for the last 30 years when thinking about antitrust and technology. I tend to group these into three separate concerns which we might want to think about in the advent of AI. The first obsession of competition economists is a phenomena called network effects. Just to remind everyone in the room, a network effect occurs when products become more valuable as more people use them. Now, network effects are a very, very popular term in competition policy. People just love saying they're a network effects. I don't know why?

There's a little bit [of a cession.] Actually, one of the reasons I wrote this paper was completely because I got irritated by lawyers labeling anything that moved as a network effect. What I would argue is if you look right now, though, it's not clear that network effects are going to operate in the same way as we've seen with previous technologies. If you think again back to Arvind's talk, do you remember he was saying, "Gosh, one of the main implications of AI in computer science or in computing could be that we end up with lots more little personalized systems." In other words, the opposite of a network effect. I think that's a far better example than I was going to use, which was if you think about the primary use case of AI right now in America, which is helping high school students cheat on English essays, there you might say, again, if you're cheating on an English essay, you really like quite a bit of divergence on output, right? There's the opposite of a network effect there. So again, this primary move of competition policy, network effects, doesn't seem to be clearly present now. We might see it as we get to the final stage, sort of a layer two being built upon our current LLM models, but it's not clear how it evolve at the moment.

Now, the next question would be one more of switching costs. Most digital antitrust cases are really about switching costs. How easy is it to switch between its search engines? How easy is it to switch between entertainment sites? This is the real question. I think there, again, we don't really have enough evidence. Again, because we don't really know what the business model is going to be, we don't really know where the switching costs are going to come from. I think probably the right forward experiment right now is to think, "Well, in the end, the most common application of our model right now is you take it into a company, you do some fine-tuning." Then there's a question about how easy it is to switch the model once you've done that fine-tuning. But again, I suspect that may be the wrong way of thinking about switching costs. But do always be thinking about switching costs when you're thinking about how AI should be regulated in terms of competition. The last thing I'll just say about that is what's wonderful about being a digital economist is that the moment you have a switching cost is the moment that everyone is innovating around getting around it.

I always used to talk about iTunes when I taught and said, "There's a brilliant switching cost. No one's ever going to want to leave their iPod when they've all their tunes on it." Then, of course, came Spotify to just make me wrong. Again, just remember the switching cost is always going to be transient. Now, in terms of data as an essential facility, there's often this idea there may be barriers to entry to do with data. But when you actually look at the training data set already out here, this is a screenshot of a website named Hugging Face, which provides million terabytes of data for free. You have to second guess it. All right, so just what did I end up saying in this paper? Or what's the big conclusion? In the end, it's really hard to be too definitive about competition policy when we don't have a clear monetization strategy. My concern right now is this our eagerness to regulate antitrust. We may end up regulating the analogy of the dial-up ISPs because we don't really understand what's going to be built upon the essential innovation. Anyway, what we do know, though, is that in some sense, the way we think about competition policy will have to change as a process, because as we just heard from Arvind, a lot of how we might think about law being conducted may change in the long run.

All right, with that, I will say thank you very much and give you over to a development economist.

Daniel Björkegren: Thanks so much for having me here. So today I wanted to talk a bit about a couple of concrete examples of projects that are on AI in low-income settings to give us a sense of what this technology is capable of and put a little bit more concreteness around some of these projections going into the future. The main project I'll present today is joint with three of my students, you know, Divya and Dominic all at Columbia. Then Paul Atherton and Oliver [Garrod] are at Fab Inc. They created the chatbot that we'll be talking about. But before we talk about that, I wanted to zoom back and think about an earlier example. I want to make the point that we may think of artificial intelligence as some altogether new animal. I think Arvind has made this point that no, it's actually an extension of a normal technology that we were familiar with. There are cases where machine learning has already impacted low-income countries, and I think it's useful to start from that premise and really understand what happened there. So one project that I've worked on is the need for financial services. There are about 1.7 billion people around the world who lack access to formal finance, which means that not only do they not have bank accounts, they also don't have credit histories. And so as mobile phones came along, people began to be able to use mobile money to transfer funds and to store funds. And then it became an issue where you could potentially imagine giving out loans via this platform, but it's very hard to know whether someone is worth the credit risk unless you have a credit score of some sort. So I suggested that we might be able to come up with credit scores just based on how a person uses their mobile phone and predict their creditworthiness. And for example, when you use a mobile phone, it records patterns of how you place calls, how you manage your balance over time, et cetera. You can use machine learning, even though if you and I, as humans, look through this set of indicators, there might be 5,000, 6,000 indicators, we would have no idea how to figure out and who is a good credit risk just based on those indicators, machine learning can connect the dots and find patterns that are not obvious to us.

In the study that we did in a South American country, we had a sample where some of the people also had credit bureau histories, and all of them had mobile phone histories. We find that when you rely on the credit bureau score, you're able to predict whether someone will repay a credit. Not perfectly, but you have some ability to do that. But as you go to people If you go to people who have less and less information in their credit bureau file, the ability of a credit bureau score to predict repayment goes down. If you go to people who have no information in their credit bureau file, you can't predict whether they'll repay a loan at all. In contrast, we found with this digital method using the behavior in just how you use your phone, we could come up with a credit score. It wasn't perfect, but it was something, and it performed about as well as the standard credit bureau scores. But the crucial thing is that we can also create those scores for people who are completely missing from the

financial system. And so machine learning really allows you to extend loans to people who were omitted from the financial system.

This has become a foundation for a new industry providing quick loans via mobile phone called digital credit. It's available around the world. And so you could imagine that we might have thought that the story would end there and that this product is out, now maybe credit is solved. I think credit answer, and this comes to some of the points that Arvind has brought up, is that there's a lot of difficulties after that point. Our team and several other teams tried to understand, was this actually improving people's lives or could it be that these were getting people into debt traps? These loans are often short in duration and have high interest rates. We found that welfare impact seemed to be small. On that, they seemed to be a little bit in the positive direction. But there's also a lot of evidence that firms who are extending loans were also taking advantage of consumers. This new industry required also some adaptations to our infrastructure for monitoring harms and regulating credit. I wanted to put this up as an example to suggest that the game doesn't end once AI enters and tries to solve a problem, that it raises a whole bunch of other questions that we need to address as a society.

Now, that is one example of what we would call predictive AI or predictive machine learning, where you typically make decisions from digital data. Credit scoring is one example. A couple other examples that have seen some promise in low income settings are one targeting digital aid, where you can use the same type of data that Josh [Blumenstock] has been doing a lot of work on, and then also weather and flood forecasting. Now, you hear you can use satellite images. One particular useful use of machine learning in this case is that often we are working in settings that have little access to data, but we may increasingly have these digital data sources, and machine learning allows us to extract signal from them to make better decisions. But the reason we're here is not just to talk about predictive AI. What people are excited about is this new era of AI where these systems are able to do more. And You probably heard this called generative AI. And this is the technology underlying chatbots that can create images and video and all of these more recent things that you've seen in the last couple of years.

So I'm going to now think about one example of what generative AI can do in this project that we were working on with teachers in Sierra Leone. But I want to first think about connectivity and information in Sub-Saharan Africa. So right now, 85% of people in there in Sub-Saharan Africa actually are within coverage of signal that they could use to access the Internet. So Africa is, to some extent, quite connected or has the capability of being quite connected. Despite that, only 37% of people are using the Internet. And if you look at how people use the Internet, even among this group of users, what you'll find is a fairly striking pattern. People use social networking, people use instant messaging, like WhatsApp, but people are much less likely to use what maybe you and I would think of as the bulk of the Internet, which is using the web and web search. Now we're going to be thinking about this problem or this gap more specifically in the population of teachers in Sierra Leone. Teachers are knowledge workers who you might think would be using the web a lot, but many of the teachers in our sample live in remote areas where Internet connectivity is quite poor.

So in our sample of teachers, 95% of them use WhatsApp every day. We find that only 3% use web search each day to look for classroom materials. And so I'm going to... Our partner organization introduced a chatbot over WhatsApp that allows teachers to access AI. And what they find is that teachers end up using this on balance a bit more often than they use search. So our paper is really trying to understand, what is going on here? What is AI doing that could be useful here? There's a couple of reasons why you might think the Internet is not doing a good job here. One is content. If you're in the US and you search for something, what you find is content from the US. If you're anywhere else in the world, when you search for something on the Internet, a lot of what you find is not local. A lot of it is foreign. So content is lacking in most countries around the world. And the second issue is bandwidth. If you ask Africans in a survey, "Are there limits to your use of the

Internet?" Most will respond, "Yes." And most will report that the cost of data is the major reason why they don't use the Internet as much as they otherwise would.

11% also report that the network is a struggle. Now, we're going to think about this in Sierra Leone, which has a GDP per capita of \$509, so a fairly low-income country. And we're going to be working with teachers who are working in schools that are fairly overloaded. And only about 6% of primary schools have electricity, and 1% have the Internet. So when we're thinking about the impacts of AI here, this isn't a setting where you're going to drop some magic tablet and it's going to magically teach all of the kids. That may come in coming years, but that's not immediately. It seems to be on the table. So in this context, the partner organization created this chatbot called the Teacher AI, and it basically links WhatsApp to GPT with a prompt that says, "Please be helpful and make useful suggestions for teachers." We find that teachers ask questions like, "Write a lesson plan to teach two-digit addition. How can I support students who fall behind? Describe immunization and give examples for a fifth-grade class in Sierra Leone." Let me just give you one example here of how the chatbot responds. So for the query activities to teach shapes, the chatbot came back and it knows the teacher is teaching grade six, and it gave some examples of creating an art collage, creating Bingo, just activities that engage the students, which seem pretty straightforward things that you might have read in a textbook, but it's tailored to this particular query.

Now, a first question is, how do teachers use this chatbot? We analyzed the queries they submitted. This is a sample of 469 teachers who asked about 40,000 queries over the span of a year and a half. We used AI itself to classify these queries into categories to understand what is it that teachers are using this for. The biggest category is concepts and facts. This is often things that you might find on Wikipedia. Though teachers also use this system for lesson planning assessment, by creating new exams and supporting their writing, and also for professional guidance. We also asked them, "When you use AI rather than search, what are the reasons?" And the top reason was that they said that AI gave them a useful answer right away and that it was concise. When we asked the reverse, not as many responded, but the modal response was that the web has unique content that AI does not have. Now, teachers seem to rate AI as quite trustworthy, but we all know that AI can make mistakes or hallucinations. We're going to come back to this later to understand why are teachers trusting this system. Then to flag, a few of these teachers also mentioned in the survey that they stopped using the web once they got access to this AI tool.

To understand why they seem to be using this tool, we took the queries that they submitted and we classified these into threads and then took the first query for each thread, which is the one that would have had the most context. We'd know how the Teacher AI would have responded, but we want to know, had that teacher taken that same question to the local search engine, what would they have seen? The most common local search engine in this case is google.com[.sl.] And so we submit that same query to Google and scrape the top five search results. And then we're going to compare what is the information that this teacher would have seen had they submitted this to a search engine on the web versus what they got from the AI response. Now, the first comparison to think about is that we came up with a ranking of how similar the AI response is to existing search engine content. We came up with a score, which we call the contained score, which ranges from 1, which means that the content that the AI generated is very new and it's different from the search responses.

Or 10, which would suggest that the content in the AI response is completely contained in one of the search results. We find that the majority of the queries that teachers submit come up with responses that overlap with existing content. It's already available on the web. They just use the chatbot. There is this small tail where the teachers are using this tool to generate new content, like a custom lesson plan or a new story, but it's a minority in this population. So they're using AI mostly to retrieve existing content. The second thing is that only 2% of the results that come back from Google are actually from Sierra Leone. And most of the results come from the United States, and

very few come from similar countries to Sierra Leone at all. So much comes from the United Kingdom or Australia, for example. And this means that not only is it foreign content, but also it is tailored for foreign users. And if we're accessing the Internet on our desktops or laptops with fast Internet connections, that content is tailored to us. And so when you look at the data required to transfer a search result, the average search result that a teacher would have seen in our sample is about 2.5 megabytes. That includes both the text, syntax, fonts, data, a whole bunch of content that's easy for us to download. The AI response was under a kilobyte. And so we find that the AI system is over 3,000 times more data-efficient than using the web. And while that data is easy and cheap for us sitting in offices, these teachers are paying per gigabyte of data. So this entails a cost. And so if you look back in time at the beginning, prior to this study, if you were to submit a thousand or download a thousand search results from this sample, you would have paid about \$2.50 in bandwidth fees. And we know that there's a discussion that AI is expensive to use. It's true that AI requires not only the bandwidth here, but it also requires a server to compute the response. If you went back in time prior to the release of ChatGPT, so earlier in 2022, an AI chatbot in this setting would have cost about \$30. It would have been about 13 times more expensive for these 1,000 queries. Over time, the compute cost for AI has declined. During most of this study, they cost about the same.

But then with recent drops, the AI chatbot has become an order of magnitude cheaper. We find that it is actually 87% cheaper to use the AI system in the setting to get a response than to access a search result. This seems to suggest that there's a lot of potential to do better by communities that have low bandwidth connections. That would be fine if it were cheaper, but we also want to make sure that it is not providing low-quality information. We also asked teachers to rate the system. We find that in these ratings, teachers rate the AI responses as more relevant and helpful. That wasn't surprising to us, but we were surprised that teachers also seem to rate the AI responses as more correct and less likely to include inaccuracies. Now, we all know that these systems make mistakes. In this sample that we tested, we found hallucinations in about 3% of the responses. So there definitely are mistakes in how this AI system interacts. But it seems like the teachers are using it for fairly basic questions for which it does a pretty good job. And also, what this exercise illuminates is that when I sit at my desk and I use a Google search, I don't expect one particular search result to be true.

I expect to have to vet multiple and triangulate between them to come to truth. And when we asked the survey respondents here to do an exercise just looking at whether a search engine result was true, they seemed to rate it less strongly. This paper, we really think of as showing that AI here can do something somewhat simple. It can reformat content to work better on small screens and [small] collections. It reveals that the information services that we're providing to much of the lower income parts of the world are not meeting their needs, and there may be ways to design better systems. This is a pretty early system. It's pretty simple, but we think it suggests that there may be a lot of different directions we could go. A few caveats. One is that the less educated teachers seem to use this system less. We do think you need some education to be able to be productive here. This is English-speaking. These models perform a lot better in English, and it seemed like these teachers needed some training. I also wanted to plug that the partner has an AI for education website that has more resources that they're doing useful things.

But anyway, very excited to think about this question about what opportunities there are to actually embed AI in low-income settings. I think there's a lot of great things we could do, but looking forward to your questions.

Gaurav Nayyar: Thank you, Arvind, Catherine, and Dan for your excellent presentations. Han Sheng, maybe I could request you to kick off the discussion and maybe Can you tell us your key takeaways from what we've just heard. Thank you.

Han Sheng Chia: Sure. Thank you so much to the three presenters for really such in-depth exploration and in some sense, this pushback to a lot of the narrative that we've been hearing

coming out of Silicon Valley, that says this technology is going to transform how the economy and society behaves tonight, tomorrow. I think, firstly, pushing back on this spirit of inevitability, I think is incredibly important. I think there's a lot of marketing that hides very sloppy thinking. I think that the authors today challenges to think about what are the specific pathways that determine whether these grand outcomes could or could not happen. Will we see high adoption because the technology is inexpensive and valuable to teachers, as Dan has pointed out, or will we see lower and slower adoption because as Arvind has pointed out, these capable technologies get intermediated through business logic, workflows in a firm that's trying to adopt this technology. It's very helpful to hear these very specific examples from both Arvind and right now. I think for many economists here, the concept of pathways and mechanisms for impact is not new. For the last 20 years, many of you in this room have probably run experiments and studies to try and tease out what exactly it is that drives the development outcomes we care about.

We've learned that sending kids textbooks is not necessarily what will drive learning outcomes, but that teaching at the right level, targeting instruction to a child's learning level is particularly effective. I think the question for us now, as we think about whether we're going to hit our development outcomes and whether AI can help us with that, is can be applied specifically to those pathways that we already know, to those mechanisms that we have already demonstrated to be particularly effective in low and middle income settings? At the end of the day, and I think Arvind has talked about this, I think it's a bit of a mirage to look at capabilities to predict outcomes, but rather we should look at specific pathways and mechanisms. It doesn't matter if our AI is outperforming every math benchmark out there. It doesn't necessarily translate into the user becoming more numerate. I think that's a fantastic contribution, I think, from this group. I do want to quickly point out that Dan has presented some empirical evidence, but really, Catherine and Arvind have made very bold, forward-looking predictions into the future. I read all of their papers in detail, and Catherine has a fantastic last line in her paper where she says, and sorry if I'm butchering the takeaway of your paper a little bit, but you say, "I think competition policy doesn't need to change that much that currently maybe our regulatory frameworks might already be quite suitable for the age of AI."

You say, "But if you're having a real knee-jerk reaction to this prediction right now, I challenge you to think about all the ways I could be wrong." The interesting question is not so much, Am I right or am I wrong? But why am I wrong? Perhaps it's a spicy question that I can I can tee up for this panel. I love maybe in the comments through the Q&A to try and play the counter argument a little bit. I'd love to hear what are some credible arguments for why your positions today could be wrong? Could superintelligence actually transform us overnight? What would a credible counterpart say to that? But let me turn it to Q&A.

Gaurav Nayar: Thanks, Han Sheng. Maybe I'll add a question from my side as well, and then maybe we can get some responses. I just want to pull a word out of Dan's presentation, the title, Leapfrogging. I think we've heard, I think time and again, that it's difficult to leapfrog general purpose technologies, although we have seen instances of leapfrogging sector-specific technologies. I think the landline and mobile phones was the most commonly cited example. I think, Dan, what you presented tells us something a little bit different in that people are actually not using web searches, but are going straight to AI searches. That's, I think, an interesting result that we see. But I think in general, just linking it to all the other presentation, so it's the promise of leapfrogging, or is it that you can get more from spending or making similar investments? If you are making investments around digital infrastructure or around smartphone access, it's just you can get a lot more from that than what you were getting before. Is that an aspect of leapfrogging that's important? Or is the claim that you're making a little bit bolder? Because it a little bit sits, I think, opposite to what Arvind said, this is just a normal technology.

Are we going to see something that speeds up the diffusion, even if it's a normal technology? Maybe if you could just take some of the questions from Han and this one, and we can keep the discussion going.

Daniel Björkegren: Great. I'll say something maybe a little bit spicy, but every technology seems normal in retrospect. I think there's been a lot of technologies that have had major impacts that were quite destabilizing at the time. I think that will be right. I don't think that diminishes what impact it could have. On the question of leapfrogging, I think maybe it's helpful to split this into two buckets. These are technologies that are seeing a lot of investment from Silicon Valley and from wealthy countries, and that will improve its usefulness for many of the tasks that happen in wealthy countries. Many of those will spill over to the equivalent tasks in low-income settings. You think about office jobs and so on, knowledge work is going to be exposed around the world. I think where a lot of my work is thinking about is, are there particular wins in low-income settings that might otherwise be overlooked? And are those going to happen automatically? So that mobile phones, for example, were just such a great technology, and the private market really allowed that to spread around the world and reach almost everyone with signal and communication technology. Or is this going to be a type of technology that is going to require some active policy to really push?

I think maybe when I focus on leapfrogging, I'm not saying that overall this is going to always cause leapfrogging, but that I'm drawn to these places where there's something interesting happening in low-income communities and debating about whether they need some policy support.

Gaurav Nayyar: Thanks, Arvind. I'm not sure if you had a... Yeah.

Arvind Narayanan: Sure. This is all as a computer scientist. I was listening to Dan's presentation, and I had two simultaneous reactions. One, admiration for the research, and the other is horror at what the web has become. If you had gone back 30 years and told people that this is what the web is going to look like in 2025, the pioneers of the web, that an average query for an educational query, not even entertainment or a video or anything, it's going to be two megabytes, they would not know how to deal with that information. That was beyond the worst case scenario of what they would have considered. All of the advantages you were talking about with AI today, those were a lot of the reasons that people were excited about the web originally. Because it's this tremendous educational tool. You can compress a lot of information into just a few bytes of text, and that it's going to have these tremendous effects, especially in a development context. I think the wounds are entirely self-inflicted. Technologically, of course, things have not gotten worse. You can still have very, very efficient web pages. We have chosen not to because search engines need to monetize and the people creating those web pages want to monetize, which goes back to Catherine's talk, which is that we're in the pre-business model stage, I think, of AI.

I think it's worth thinking about 10 years from now, what are those business models going to look like? Could they make AI look a lot like what the web looks like today? I don't necessarily want to constantly be the pessimist, but it is something I think about a lot when I look at the claims around Enshittification, for instance, if folks have heard that term, the way that social media platforms and other two-sided marketplaces have the rent-seeking effects of degrading product quality in the interest of making things more addictive and making more revenue. Again, while completely agreeing with Catherine that it's very early days, I think it's worth thinking ahead about what are some of those potential negative effects that are coming to these new platforms and what perhaps can be done from a policy perspective without being too prescriptive, without jumping the gun on competition policy to try decrease the probability of some of those future negative consequences.

Gaurav Nayyar: Catherine.

Catherine Tucker: All right. So I love the question, which was, "Tell us why you're wrong." Such a wonderful question. So I've been thinking about ways I am wrong. The first thing I would like to highlight about why economists are often wrong is this analogy. Our problem is that often when we

think about labor displacement. We have in mind something analogous to thinking that what Uber is doing or ride sharing is doing is replacing taxi dispatchers. And that, in some sense, is completely underestimating the real change. If you think what ride sharing has done is it's completely transformed the process surrounding how you get into something akin to a taxi. And that's what makes us wrong, and potentially, maybe, that could be what makes Arvind and I wrong, that we just don't anticipate enough how much process change there will be. The second reason I think we could be wrong, or I don't know, I meant to be only talking about myself, and now I'm making everyone else wrong. But I think in a talk like this, it is very easy to fall into the trap of calling everything artificial intelligence, when often what we're really talking about is something more akin to machine learning, predictive analytics.

This reminds me of, and again, that could lead us to make mistakes. I'm always reminded of one of my more famous papers in the realms of AI. I was looking at how machine learning did, these are machine learning algorithms run by large credit bureaus and similar data brokers. What I showed there was that machine learning can generally, if given, web search or web browsing data, classify gender right about 50% of the time. If you think about it, that's quite funny. Given that machine learning can often get things wrong, predictive analytics, as of now, can often get things wrong, we're been being very bold by overselling it in terms of this AI label, which I think we've all been using. I think those are the ways I'm wrong, right? Misusing language under predicting process change.

Gaurav Nayyar: Thanks, Catherine. Maybe I have one more question before we open up the floor to questions. I wanted to pick up, I think, one of the pictures that, Arvind, you had early on in your presentation, just on invention, innovation, and wider adoption or diffusion. I think just because this conference is about development, I just wanted to see if how we can place or contextualize developing countries along this process. Perhaps there is room for them to be in all parts of the process, but more so in one particular part of the process. Maybe Catherine, also extending your talk about competition and value. I think the sense that people have is that all the value is in the invention stage that Arvind put up, or perhaps a little bit in the innovation, but not so much at the bottom end. How should we think about that?

Arvind Narayanan: Yeah, I'll maybe say a couple of words, and then would love to hear from folks who have more expertise in developmental economics. I'll start by saying that developing countries focusing industrial policy on data centers, for instance, seems bonkers to me. I'm very convinced by Jeff Dings' book, *Technology and the Rise of Great Power's Looking Historically on General Purpose Technologies*. It ended up not mattering that much in which country the innovation happened. What was much more useful to explain the variants and outcomes between countries was how good those countries were at diffusing those technologies throughout various sectors of the economy. Policy can play a big role there. Just looking at regulation as one form of policy can both inhibit diffusion and spur diffusion. My favorite example of this, really trivial, but I think really powerful, is in the early days of the web, e-commerce was just such a strange notion to people. What actually helped that along in the US, but I think it applies anywhere, is the East Side Act that said that companies are forced to accept an online signature in any contract in which they would have accepted a physical signature. This one thing really changed the game because it allowed people to digitize their processes without worrying that suddenly they're going to run into a roadblock where a party they're trying to contract with will not accept online ways of doing things.

This was a huge project. It's a simple example of how policy can play a big role in diffusion. That's where I think the impacts are really going to be felt in terms of where countries can leapfrog, et cetera, as opposed to putting up a data center. It's not clear what they're achieving by having the model training done in their country when it's owned by an American company anyway and then can be used by anyone in the world. So something does not compute as far as I'm concerned.

Gaurav Nayyar: Catherine, Dan.

Catherine Tucker: Well, I'll just say, you may, again, you may be thinking, sitting here looking at me, going, "Of course, it's a business school professor going on and on and on about monetisation models." So predictable, completely irrelevant for my world. Now, I would like you to, though, have this forward experiment about why monetisation models are important. And if we were to go back in time to 1995, and I was to tell you that monetisation models would develop, which meant that mapping would be free for the rest of the world, you would have thought me completely insane. And especially if I told you that that was going to be advertising. Again, sounds nuts. And so I would argue that in some sense, where we land with monetisation models is going to be important really just in terms of access, because it's going to very much shape access internationally.

Daniel Björkegren: Yeah, so I agree. A lot of the work here is figuring out what these What these systems are capable of and how that integrates into local economies. What does that mean for business processes like your specific business process? There is a challenge here in that some of these models will already slot in well into some tasks. Some of them will currently fail for reasons that are locally specific. It may be bad in the local language. It may not know the local conditions. I think there's some danger here. I think we need a combination of exploration by entrepreneurs and et cetera, really pushing it to understand how this technology can be useful in local conditions. But then we will, at some point, need more coordinated investment to make sure that if there's latent capabilities where these models could be useful for many people around the world, but they would require some public intervention, how do we design systems for that? I think, are the business models going to support that in wealthy countries where this will happen automatically, or do we need other institutions to step in? It's a huge question.

Han Sheng Chia: If I may just build on that and to Arvind's question about why are our lower income countries just so obsessed about building data centers and having their own model development capabilities, I think the pushback to that point I've been hearing is that competition among the major AI labs in high-income countries are not going to deliver on the services that ultimately we need in our local context. I think some of the obvious cases here would be minority language development. There's absolutely no business incentive, very difficult business models that would incentivize the major tech labs to come in to do that. Competition policy, to Catherine's point, focusing on lower prices for consumers, allowing switching costs. Again, that doesn't quite account for that. I think just what I'm hearing from policymakers in developing countries is that's the push for the need for local capacity.

Gaurav Nayyar: Thank you. The floor is open. Please come up to the mic. Short introduction and a short question would be ideal.

Daria Taglioni: Hello, Daria Taglioni, World Bank. So thanks. I want to, I don't know, I guess, reflect on two takeaways from this chat. The one is that AI automates routine tasks and there is nothing new. And the second one is that we are in the pre-business model. It might be true that we are in the pre-business model, but I think it's important to think that this pre-business model is already reshaping labor and it is having systemic implications. It has to do with the fact that AI is based on the gig economy. The supply chain of AI is structured around squeezing the maximum value from vulnerable populations. The annotators, which are the people checking the code, are paid 90 cents per hour. They are taken from countries that normally are imploding, Venezuela, Kenya, and then in turn, all the countries that were imploding after the COVID crisis. And so all these people on top, there are algorithms that check pics of productivity, and they take away tasks from people or communities that are becoming solely reliant on that gig economy because everything else has imploded. And so all these people are part of a growing global informal, unregulated workforce where there is some basic laws around labor that we established some 100 years ago that are not reflected in the gig economy because there is no stable employment, no benefits, no protection, and not even an employer you can go to and complain.

Hence, my plea is and comes to some of the points that were made, is that we need to recognize in this pre-business model and quantify and price in the human and social costs. Otherwise, the conversation about AI will be dangerously incomplete. So labor protection, transparency, fair compensation, community control is necessary unless we want to replicate some colonial ear extraction and call it progress. Thanks.

Gaurav Nayyar: Thanks. We'll just, in the interest of time, just collect all the questions and then we can have one round of replies. Thank you. Please, no commentary, just a short question.

Ilaida: Okay. Hi, this is Ilaida from Syracuse University. Thank all for your thought-perfecting talks. I specifically appreciated Arvind's contextualisation of AI within the history of technology innovations. Power and Progress from Acemoglu and Johnson received a lot of praise from scholars from a variety of disciplines as they pointed out to the fact that the fusion of tech innovations is not a natural evolution, rather what will be automated and how is political decisions and that institutions are central to that decision making. My question is, what do you think about the role of development organizations in giving policy direction and AI policy making? Thank you.

Audience Member 3: Thank you very much for your articles. My name is [Unintelligible], and I'm a founder of an AI startup. My question for the three presenters, I assume that the work that you have done was based on the generative AI that exists today. There's a lot of indication that agentic AI would be actually the most important when it comes to the impact on economies and societies. Do you think that if you look at your work again with the prospect of agentic AI, do you predict the same results or not? Thanks.

Daniel: Thank you. This is not really a fair question, but I still want to hear your insights. Daniel, Mexico has a huge education system crisis right now. How do you see the use of AI for education development on the serving areas of the developing world, especially when infrastructure challenges are acute and government has stepped back from education as a policy priority? Thank you.

Abdulla: Hello, I'm Abdulla. My question is broad, but also specifically for Catherine as well. Our role as economists is to learn from the mistakes of the past that we've made in terms of global economy and not make them in the future. And Catherine mentioned that economics doesn't change, but if all of this predictive inference is wrong, and we've seen that they've been proven wrong, but the radiology and all of those examples. Perhaps then economics needs to change in order for our predictions to be right because the consequences of being wrong are that you might just end up in a bubble, like the dot-com bubble. The consequences are multi-tiered in that case. How do you guys think, in your knowledgeable opinions, that we need to change this predictive inference to be more right and make sure that we don't end up in another global crisis where people are impacted, technology is impacted, and this just ends up being a repetition of the same bubble again. Thank you.

Gaurav Nayyar: Thank you. Last question, and then we'll get responses.

Audience Member 6: First of all, thank you very much for the insightful discussion and the presentation. I really appreciated it. Going right to the question. It's a question regarding the use of AI in the development field as one of, if not the most versatile and volatile element of the digital public goods. So for digital ODA, in the terms of the soft infrastructure, what would be the main incentives for donor countries to actually provide aid to the recipients? And on that note, what policy implications can we infer from that? Thank you.

Gaurav Nayyar: There's one more. Please go ahead. But we'll have to cut it there because...

Maya: Thank you so much. My name is Maya. I really appreciated the fact that you were challenging this view that AI will replace jobs and instead will change them.

So my question is if you envision this idea that maybe AI will create new industries on jobs that will support developing countries.

And if so, what would you picture that those new jobs or industries to be like?

Gaurav Nayyar: So maybe you could take a minute each to pick and choose whichever question you want to respond to. Thank you.

Daniel Björkegren: Great. Yeah, lots of great questions. I'd say in terms of... There are a few questions that were thinking about the future, new jobs, et cetera. I think in the near term, there's some wins where AI seems quite amenable to improving education and health service delivery in many countries that can't afford great human providers for all people. I think there is a lot of opportunity. It's difficult to know exactly how to integrate AI assistance into classrooms, and I think that's a challenging problem. We don't quite know what the optimal structure of teaching is that [unintelligible] this. There's a lot of questions there. Then there also is a lot of uncertainty about the long-run trajectory, about what capabilities future AI will have, and also as it's increasingly integrated in the economy, if these bold predictions about economic effects will pan out. And I think there's some chance that it will. I think the job implications for low income countries, I just have a lot of uncertainty about what that could look like. I do think we will be surprised by the new jobs that emerge, though.

Gaurav Nayyar: Catherine.

Catherine Tucker: I just wanted to highlight two things I took from the questions about what we're probably getting wrong right now in the economics of AI with our focus on high income countries. I think the first thing is that, as you can see here, we basically hate going off the task-based model, saying thinking about AI as task replacement is not helpful. But I thought what was useful is just a reminder that the West or developing country or high-income countries have outsourced a lot of tasks to low-income countries. Therefore, we should think a little bit about that before we just say tasks will change and shift and so on. The other thing I want to say is, again, the economics of technology tends to draw a lot of reassurance from the fact that when we study what happened to telephone operators with automatic switchboards, when we look at what happened to accountants, when the calculator was invented, it looks really good. In 10 years. But we tend to forget about the pain, the upheaval, the horribleness of the 10 years in the intervening period. I think that's, again, something that we don't do enough of.

Arvind Narayanan: Thank you. I'm going to address briefly two of the questions. One was about generative versus agentic AI. Our framework is not specific to generative AI. In fact, I lead a team of technical researchers. The main area we've been working on over the last couple of years is agentic AI. That's been deeply on our mind when we wrote AI as normal technology, and it is agnostic to what form of AI will be at the frontier in 10 years, 20 years. I want to give one note of caution around the hype around generative AI. I say this as someone who finds it very useful. I spend several hours a day using coding agents, and they're fantastic. But when you look at the agentic AI that you would need to have rapid transformation of the economy, I think you need three things. One, you need them to operate with little human intervention, relatively autonomously. If you have a human in the loop, you're not really achieving that much in terms of productivity gain. You need them to operate in tasks which have high costs of errors. If you're just getting responses from a chatbot, that's very different from doing medical diagnosis or something like that.

We identified these three aspects of autonomy in agency in generality. It has to be one general purpose system as opposed to something that you had to spend 10 years developing for a particular worker. When you look at these three factors, there are zero deployed systems right now, as far as I could tell. We predict that it's going to be zero for a few years to come. Unless that changes, while agentic AI is very exciting from a technological perspective, I would not count on it being a counter example to the kinds of things I talked about. Okay, so very quickly on new jobs, we're big on the potential for AI to create lots of new jobs. I obviously can't predict with any certainty, but I think policy needs to be very attuned to this possibility. I really like David Autor's paper called Rebuilding Middle Class Jobs Using AI. He talks about historical transformations that went from the artisanal

expertise of, let's say, potters to the mass expertise of people who could use machines to perform those in a much more efficient manner. That was a lot more people. Depending on demand effects, you might actually have an increase in labor if the demand for pottery goes way up because it's cheaper.

A similar thing might happen with AI. You might have a result of pervasive cognitive automation. Could be that you have mass expertise of things that today require 10 years of training like law and medicine. A lot of the aspects of those professions could be performed by people with, let's say, six months to a year of training who are using it in a way it's a human AI team. To take advantage of that, what we will need is new rapid credentialing so that you can have some licensing of these people, but it doesn't involve the same 10-year pipelines that we have today.

Gaurav Nayyar: Thanks. We've passed our time. I will not try to summarize the session, but the three big takeaways for me were, one, and I made a note of these, one, be humble and acknowledge the uncertainty. Two, the more things change, the more they stay the same. Three, be serious about your research. Collect the evidence to help guide the decision-making, which, as it turns out, I guess, are just life lessons. I would like to thank the panelists for their wisdom, and I thank you all for joining. Thank you.

Karolina Ordon: Thank you so much. We're going to have a 10-minute break. Coffee is served on this side. When we come back at 4:20, we're going to have a session about one of the most urgent issues for developing countries, which is pollution. So please be back on time. Thank you.

[END OF TRANSCRIPT]