

THE BROOKINGS INSTITUTION

FALK AUDITORIUM

CLOSING THE DATA GAP FOR AI POLICY: LESSONS FROM THE STANFORD AI INDEX

Washington, D.C.

Wednesday, April 15, 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 INTRODUCTION:

ELHAM TABASSI

Senior Fellow and Director, Artificial Intelligence and Emerging Technology Initiative,
The Brookings Institution

REPORT PRESENTATION

SHA SAJADIEH

AI Index Lead, Stanford Institute for Human-Centered Artificial Intelligence (HAI)

CONVERSATION ON ECONOMIC & WORKFORCE DATA

RAY PERRAULT

Co-Chair, AI Index Steering Committee & Affiliate Fellow, Stanford Institute for Human-Centered Artificial Intelligence (HAI), Distinguished Computer Scientist, SRI International

SHARAT RAGHAVAN

Director of Data Science & Director, Economic Graph Research Institute, LinkedIn

DAVID WESSEL

Senior Fellow and Director, The Hutchins Center on Fiscal and Monetary Policy, Brookings

PANEL DISCUSSION

MODERATOR: SANJAY PATNAIK

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

YOLANDA GIL

Professor of Computer Science and Spatial Sciences & Principal Scientist, USC Information Sciences Institute, University of Southern California, Co-Chair, AI Index Steering Committee & Affiliate Fellow, Stanford Institute for Human-Centered Artificial Intelligence (HAI)

ELHAM TABASSI

Senior Fellow and Director, Artificial Intelligence and Emerging Technology Initiative,
The Brookings Institution

RUSSELL C. WALD

Executive Director, Stanford Institute for Human-Centered Artificial Intelligence (HAI)

* * * * *

BELLE: If everyone can take your seats, we are ready to get started.

Good afternoon, everyone, and welcome to the Brookings Institution. My name is Derek Belle. I'm the associate director for the AI and Emerging Tech Initiative here and I will be your mc today. AIET is an institution-wide effort at Brookings, which leverages the breadth of our expertise to drive research, collaboration and institutional innovation to shape AI policy at this critical moment.

We have a great program for you today, so I'm gonna keep these housekeeping remarks very, very short. We're joined by colleagues from the Stanford Institute for Human Centered AI or Stanford HAI who I wanna congratulate on leasing a mammoth AI index earlier this week. In this session, we will see a presentation of that report followed by a conversation on the economic and workforce impacts of AI. And we'll finally close with an expert panel that dives deep on a few of the most consequential policy areas facing us today. Each of the interactive discussions will be followed by an audience q and a, during which we'll have mic runners in the audience. So raise your hands and one of 'em will come find you. We're on a tight schedule, so please keep your question short. If you're watching online, you can email questions to events@brookings.edu.

After we conclude, there will be time for those here to stick around in the room, enjoy some refreshments, and meet some of the other people that we've assembled here today. And before we get into any of that, I wanna introduce our first speaker, my boss and the director of the Brookings AIET initiative, Elham Tabassi. She also sits on the steering committee of the AI index, so Elham, please join us on the stage.

TABASSI: Thank you, Derek. I'm not gonna make a speech because we wanna dig into the, to the report, but I wanna welcome all of you to here, to these conversations today. We are extremely pleased to have the the people who led the development of the AI Index today. My name is Elham Tabassi. I lead the Brookings AI and Emerging Tech initiative where we focus on bridging the technical and policy communities and conversations and discussions, and also grounding governance frameworks in empirical research because I think as the report is gonna show the cannot or hardly can manage what we cannot measure.

So we are here today to dig into the finding of the latest Stanford AI Index report. As you all know, it's a comprehensive annual snapshot of where artificial stands globally. And rather than just presenting the data and going through the text of the report, we have the, extreme pressure of having the people who led the development of the report, the people who collected the evidence, made the editorial call, put the report together here in the room with us. So, I'll, I'll be short and get into the conversations about the numbers, the surprises, and the questions that are left to talk about in the next 80 minutes or so.

With that, we are gonna start with a presentation on the report, and it is my pleasure to introduce the reports lead, author Sha Sajadieh. Sha joins Stanford HAI after spending over a decade in product leadership roles at leading tech companies. She serves as AI Index lead at Stanford Institute for Human-Centered AI or HAI. And manage, manage the direction and development of this year's report. Congratulations on the report and please join me on stage for your presentation.

SAJADIEH: Hello everyone. Good afternoon. Thank you Brookings and the AI and Emerging Technology Initiative for hosting this conversation. As Elham mentioned, I'm Sha, the lead author for the AI Index. And before I dive in, I wanted to -- should I use this clicker by the way? Let's see. Here we go.

Okay, great. Great. This is, I wanted to acknowledge that the AI index is a work of a large team and a steering committee of leading researchers and practitioners across academia, industry and policy. You'll hear from several of them throughout this program including our co-chairs, Ray Perrault and Yolanda Gil.

Their guidance really shapes every edition of this report alongside our supporting partners and the data vendors that we work with to do the analysis. So the index is in its ninth year, is the ninth edition, and we cover nine chapters. For the first time we have standalone chapters in for AI and science and AI in medicine, reflecting how far the technology has moved into specialized domains our mission hasn't changed. We bring rigorous independently source data to a field that is moving quite fast.

And really the thorough line of this report is that AI is, is scaling faster than the systems built to absorb it. So governance frameworks evaluation methods education systems, the data infrastructure underlying it this gap runs through every single chapter of this report. And it's the empirical condition under which all of us, the policy community, and others are operating under. And so for the next few minutes, I wanna walk you through what the data shows, but also what the data, what data we don't have and both really matter.

So adoption of AI has been unprecedented. And you all probably have experienced this in using these AI tools more and more. 88% of organizations have now adopted this globally, AI and at least one business function. Generative AI has reached population level adoption within three years. It's faster than the personal computer or the internet and global corporate investment dollars that's doubled year over year into AI companies. It, despite all of this though, a AI adoption is uneven. You know, for example, the United States, we lead an AI development and investment. But at a, in terms of adoption, we sit at 28th or 24th amongst other countries with an adoption rate of 28%.

We can compare that to Singapore or the United Arab Emirates, which are over 50%. So, you know, the leaders in AI development are not always those, their populations aren't always the ones that are using it, the heaviest. And for the United States, the gap between develop and adoption really matters because we, the, the United States really leads on a lot of the inputs within the AI development pipeline.

In terms of global private investment dollars this year alone, it reached around \$286 billion, which is almost 160% year, year over year growth since last year. And that's you know, by quite a large margin in in terms of the next country, which is China. The United States also led in entrepreneurial activity, so around 2000 new companies this year in 2025.

More than 10 times the next leading country as well. But the talent picture is really shifting in a different direction. The number of AI researchers and developers that are moving to the United States has dropped about 89% since 2017. And 80% of that decline has been in the last year alone. And we believe that this is an indicator definitely worth watching.

At the technical frontier, leading models are now really indistinguishable from one another. A year ago, the top four frontier models had a gap of about 97 points amongst them. Today they are separated by fewer than 25. The US China gap was sizable in 2023, but that has nearly closed with the leading top model from the US leading by less than three percentage points. And the tools to measure the capabilities of these AI tools is really the, the models are really outperforming them. It, it becomes harder and harder to tell the difference between all these top models. So cost latency, reliability, and governance frameworks become even more important as we think about it.

But ultimately openness is declining. Over 80% of the notable models that were released by industry they were released without training code. As you can see, the open source models, which is small sliver here at the bottom in 2025 and parameter counts and dataset sizes are no longer reported for several of the most resource intensive systems.

So this is a true structural problem. Benchmarks are saturating. A test designed to be really hard for these models aren't lasting that many years before, or even months before the, the models themselves are able to outperform them. So, you know, again, the data we're having about these models and how to govern them is declining, but the systems are getting more and more powerful.

As I mentioned, we have chapters on science and medicine this year. Really exciting things and breakthroughs that are happening in both fields. AI has shifted from accelerating scientists on individual research steps to moving into full end-to-end workflows In medicine, we're seeing AI being deployed in hospital systems.

For example, ambient AI scribes and really helping pro prevent clinician burnout, which is really exciting. But there is an evaluation gap here too, because a lot of the clinical AI studies happening today are not using real patient data.

Similarly in education in classrooms, students are moving far ahead of any policy that can keep up with them. Over 80% of students are using generative AI tools for school related tasks. But less than half of schools actually have policies in relation to these tools. Teachers, you know, only 6% of teachers really feel prepared to handle the guidance of these tools.

And ultimately the way that we measure education and AI education, specifically, the data is lagging in incomplete in our best proxy is computer science at, at the moment.

In terms of economic indicators, we know the value is there. Consumer value from generative AI tools is rising, reaching about 172 billion this year in early 2026. But the labor market signals are a bit mixed. Productivity is gains are highly context dependent. There's a ton of micro studies going on and has been for the past few years that tell differing stories.

And we also do see some level of decline in employment levels for software engineers. Age 22 to 25. But again the more longitudinal data we have around this, the better, the clearer the story that we can tell. Finally, governments are acting on ai, but not all in the same way. As we know, the US is shifting towards deregulation.

There's a federal preemption of state laws as we move to a more federal cohesive framework. And what is exciting though is that more than half of new national AI strategies came from countries that previously have not been in the AI policy making space at all. And AI sovereignty is emerging as a really important principle, though still to this day.

I would say there is a gap between what that framework really means, and if there's a shared taxonomy or definition around it. And I think, you know, we all may feel some type of way about AI in this room, but, and the sentiment really does depend on where you are. So the most optimistic countries places in Southeast Asia and China, the over 80% of the population feels really, really optimistic and excited about what AI is going to do for them in the next three to five years.

And they also trust their governments to regulate it. In the United States, we move in the opposite direction. There's a high level of skepticism. Only 33% believe that AI will improve their jobs or the job market in their in their state. And they have the lowest level of trust in US government to regulate ai.

About 31% compared to these other countries that have such a high level of trust. So the report is out. Take a look. And ultimately I'm really excited for you all to hear from some of the experts. And as I mentioned the data that we have is just as important as the data that we don't have and are trying to understand more about in the field that is moving this fast.

So thank you for having us. And I will hand it over to Derek. Yes.

BELLE: All right. Thank you, Sha, for that insightful presentation and giving a survey of some of the big themes that were identified in the report. Our next two segments are really gonna sort of tighten the focus on some of the most significant policy challenges that face us today. And we're gonna begin with its economic and work workforce impacts and the empirical basis for understanding them.

So we have a panel of a, a few esteemed experts and I will pass it over to Dave Wessel, who is the director of the Hutchins Center on Fiscal and Monetary Policy, and a senior fellow in our Economic Studies program here, Brookings. Take it away, Dave. Thanks.

WESSEL: Good. Well thank you very much and in case anyone, I am definitely not one of the esteemed experts, so that'll identify who else is on the panel. It's really exciting to have Ray Pinot here, Perrault here, who was one of the co-chairs of the AI index steering committee who saw his picture up earlier. He's a distinguished computer scientist at SRA International, a big nonprofit research institute. And we were talking earlier, I love this line. In his bio it says from 2002 to 2009, he was co-principal investigator of the Calo project, a large multi-institutional DARPA funded projects whose objective was to build an intelligent office assistant that learns through interaction with its user and the world, a product that has not yet arrived on my desk, but is actually one of the seeds of Apple Siri.

And Sharat Raghavan is director of data science and research at LinkedIn's Economic Graph Research Institute that's using LinkedIn data to analyze labor markets and make policy recommendations. So we have a very short panel, 15 minutes, and then we'll have five minutes for questions. I do wanna say that this report is an amazing compendium of really interesting information and little tidbits, and someone did a lot of work to give you the highlights of every chapter, which unlike most of the chapters, when you read these highlights, they tell you, like on page three you can learn this, it's actually synthesized very well.

I suspect ChatGPT was involved, ChatGPT was involved. So one of the themes of today's conversation is data at at, as has said, in the absence of it, so. If we're thinking about the impact that AI has on the economy and on the labor market in particular, what are the challenges we face in collecting this data, which is important if we're to come up with good policies.

What are the obstacles? Do you wanna start? Right?

PERRAULT: Thank you. It's a great to be here. It's first time in this organization, and then a real privilege. So I've been with the AI index since round one. The first report we published in 2017, I think was 60 pages and it kind of grew after that.

And there are a lot of themes that have been constant throughout the iterations of the report, although every one of them contains contains something new. One of the bits that so, so I've had a, a, a number of kind of dista for data that I wish we had and that we don't. And one of the.

There, there've been two kind of big overall directions that we wanted. One was to present data in a longitudinal fashion. So if you come back year after year and you find out what's changed and that's by and large what we, what we manage to do. Every so often there's a bit of a chapter in which we only have one data point 'cause we think this is gonna be important, but that's all we can tell you.

And then the other one is geographic distribution. We really would love all this data to be equally valid across the world. And those are both, as you know, all of you know, big challenges. And one of the geographically oriented ones that I wish we could do better on is an assessment of the size of the.

Manpower in each country that is competent in, in ai. And we do that in bits and pieces, in different in different countries, but we have no uniform way of bringing it of bringing it all together. And so this is something, if any of, you know, of, you know, ways of thinking about this. And, and my sense is that size of the competent man of the, the competent workforce in the area is more important than.

What gets invested in companies? 'cause mostly what gets invested or numbers of publications in in universities or whatever, are a reflection of the number of people who are working on the problem. And there are big differences when you look at investment, for example, between how investment is spent, public investment is spent in different in different parts of the world.

So in the US you apply to NSF you get a grant. Most of your grant goes to pay for graduate students. Some of it may pay for your summer salary. You go to Europe and their salaries are 12 months. So you couldn't compare grants in Europe to grants in the US 'cause they don't pay for the same thing. So we could get, dollar numbers or euro numbers on a lot of these things. But we would have to be able to do some pretty creative translation of what these numbers mean from one environment to to the other. And then there are places like China where we have very little of it, but there are some things we can estimate.

Like we know roughly how many publications Chinese researchers put out every year. And publications are, I think, a pretty good measure of the size of the research community. 'cause you can only publish so many papers in a year. So there may be ways of taking a bunch of these indicators and bringing them together to give us something that's more comprehensive.

WESSEL: Right. When you talk about the number of people that are competent in ai, I thought you were talking about. The rank and file workforce, but you mostly seem to be talking about the research workforce.

PERRAULT: I'm, I'm talking about the, the, the workforce that is the source of the development of new new technology within the, within the country.

Right. And I think when people, I mean, that's one of the, I think clearly important parts of how people compare countries, are they gonna push it?

WESSEL: So Sharat, how much do we know about how people are using AI in the workforce? And what is it that you wish we knew more of? And how can we get it?

RAGHAVAN: Absolutely. So first of all, thank you for having me.

Just as a side note, you know, a lot of discussions sometimes about the harms of ai. So I walked into a door this morning as I was looking at a report on my phone, so don't everyone, you know, the old lesson of don't just stare at your phone as you're walking around anyways, had to get this patched up here.

Hopefully I'm at a good angle here. So

WESSEL: and was that due to AI or is that incident?

RAGHAVAN: I'm gonna blame AI on that one actually, I think that one's AI's fault. So I lead economic research at LinkedIn and you know, we have this data set we call the economic graph, which is a real time representation of the workforce.

We're so grateful to partner with Stanford and other institutions in sharing this data. And you know, one of the things we've been tracking for a while is the growth in AI skills. And we've seen AI literacy skills more than double in the past year. For members on our platform, we've seen a 70% increase in job postings asking for AI skills.

And so to get to your questions about what data would I want to have, or what data do I think we need now, we see a lot of data on exposure, on skills, on adoption. But just because AI can do a task doesn't necessarily mean it's more productive than someone actually doing it. A job is just not, you know, it's not even a collection of tasks, it's how they interact with each other.

Really what makes work fundamentally human. And so one of the things that we've been really thinking about is not only how we measure adoption, but how we think about productivity. And how AI is going to impact that in various ways. And you know, that's something that we're really laser focused on, is really trying to understand that question.

WESSEL: So Ray, if, if you can, what, we know something, we learned something from the report. As Sean mentioned, there seems to be a decline in employment among young software developers. There's like lots of micro studies that suggest this has great potential to increase productivity, although you have some interesting stuff in the report and saying in some places maybe it hasn't.

We know that you say in the report that one third of organizations expect AI to reduce their workforce in the coming year, but large scale job losses haven't shown up, and almost half the organizations don't see any change. So give me a snapshot. How much is AI affecting

the workforce today? And where do you think we can look to see where the most net what comes next?

PERRAULT: Yeah, so we're certainly seeing it being more difficult for recent graduates to find jobs. I think that's pretty clear. And, we are seeing some companies lay off people and blaming it on ai whether AI is the real culprit I think has been challenged. And in part it's, there's a claim that a lot of companies overhired in the pandemic and they're still trying to come back to something that they can afford that they would rather spend their capital on buying hardware and data centers than on, than on people.

It's hard to tell what's really what's really going on there. So, I, I really look forward to you know, real expertise on this. As we say in the report, we, the, the, the, the the direction, the fingers point in different directions. I don't think any of what we. Present is contradicts any report, contradicts any other, but they're subtly different.

WESSEL: Well, it seems to me that the report does point in the, it's highly likely to increase productivity over time. And the challenge is to diagnose how quickly that's gonna happen, how quickly it'll diffuse, because that makes a big difference on the labor market. Dislocations. Yeah.

RAGHAVAN: Yeah.

Yeah, so I think that's the question that we are constantly thinking about as well, is trying to understand where we're seeing AI's effect on the labor market.

We just published a report actually in January where I think if you look at the takeaway, if you organize, let's say, occupations by occupations that are more exposed and less exposed to ai the rate of hiring on LinkedIn that we measured by our LinkedIn hiring rate, you know, they're following in a very similar pattern.

It. And so when we look at that, we say, well, it's hard to disentangle what's happening because of macro you know, macro factors. It's hard to disentangle ai. It's also hard to disentangle the uncertainty of what AI can do, which is, I think, a different concept, right? It's not necessarily AI is affecting it, it's more like the UNC.

Is causing some people to dial back hiring. But however, if you look at like micro level studies, you can see certain occupations, like copywriter is one, for example, some occupations in legal, especially on the early career side where you do see some weakness, it does seem to correlate to increased exposure.

But again, it's very hard to make definitive statements and even on software engineering as being sort of like this you know, a, a signal of what we might be seeing. You know, the many of the large tech firms actually in last year added software engineers to the headcount. And so ironically for a lot of firms, AI is making it a little bit actually easier to produce things.

And so when it makes it more easier to produce things, you actually might increase demand paradoxically. And so we might be seeing a little bit of that. Hmm. So it's a, you know, it's a little bit hazy.

WESSEL: Great.

RAGHAVAN: Let me ask you one, can I ask

PERRAULT: one, add one bit to the software engineering one

WESSEL: bit? Yes.

PERRAULT: One bit.

There was this study published about a year ago from Meta that claimed that highly experienced software engineers were seeing a decrease in productivity. More recently meta, meta wanted to replicate that study and they didn't do it because apparently they could not get a sizable enough cohort of programmers to take the test without access to coding tools.

Right. So,

WESSEL: yeah. Right. Let me ask you one final question before I turn to the audience. You showed this sort of scary chart about the ebbs and flows of people coming to the us. You've been in the computer science business for a long time. To what extent do you worry about. Immigrants coming to the US to be in these fields and leading to leading innovation, how, how big a threat do you think that is to our future?

PERRAULT: Well, I mean, the history I think has been that immigration has been good in technical areas, let alone others. But I mean, in technical areas in the US both Yolanda and I are immigrants. And it just seems to be we're shooting ourselves in the foot by turning down, you know, brilliant people from around the world.

Look at other people around the room. So yeah, we, we need to fix this.

WESSEL: Hmm. So, public opinion data shows, as we saw in the chart at the beginning of the presentation, that a lot of Americans don't like ai. They're afraid it's gonna eliminate jobs, and they think that we don't know how to regulate it.

I mean, we clearly have a generic trust in government and institutions and elites, of which this is one manifestation, but I wonder if you, either of you have any ideas about how can policy makers best address these concerns to give people more confidence that like we can embrace this technology and we can handle it.

And to the extent it's disruptive, we will find ways to ease the pain. If you do it well, you can run for Congress

RAGHAVAN: so. Definitely a question that comes up quite a bit. You know, I think from the private sector, at least I can speak from LinkedIn and Microsoft, is that, you know, we pretty much started very early on with a set of responsible AI principles around how we use ai, the fairness behind it you know, making sure that we deploy these models for the good of the product.

And so if you think about what LinkedIn is, you know, it's to generate economic opportunity for every member of the workforce. So one of the things that we've been using, for example,

AI, is on the recruiting side, is to open it up to a skills first hiring perspective. And, you know, there's. For example, a lot of, you know, you know, inbuilt, sort of institutional you know, baggage re regarding where you went to school or where you live, where your occupations.

And so we've been really focused on kind of driving, for example, to a lot of the recruiting tools to focus on skills. So to maybe get back to your answer question about regulation, I think it's a pretty fast and evolving technology. Obviously there's a lot of like geopolitical concerns. I do think in an ideal world there'd be, you know, the major bodies of the world would come together on a line on some principles.

But as we saw in the report, it's quite a bit of a race from a geopolitical standpoint between the US and other countries. So it's how you do this in a really intelligent way. So a little bit sort of beyond sort of what we look at on the data science team, but I think it's super important.

PERRAULT: The, the confidence I think of the general public in using these tools is largely derived from what benefit they get and whether they trust the outcomes. I think. By and large, the current commercially available systems when used by most people produce high quality results that are verifiable.

And you know, it's kind of like in the early days of the internet, you can't trust everything it tells you and you're tempted to trust.

WESSEL: But in the later days of the internet, you can,

PERRAULT: but, but you learn, you learn to deal with it. I think that's a part of the, of education is to

WESSEL: learn

PERRAULT: to deal

WESSEL: with it.

Yeah. But I think that there's a, I think there are two different things going on here. One is anxiety about are my kids gonna be able to get a job? And that's one level. And then the other is all this stuff about the scary stuff of people falling in love with an AI bot and people committing suicide.

It's gonna be really hard to regulate, but it's really frightens people in a way that the electric typewriter didn't.

PATNAIK: For sure.

WESSEL: Alright. I could keep going, but I'm not going to, so I'm gonna take like three questions and they're gonna be short or I'll cut you off and then we'll let them answer and we'll move on.

In the back. So stand up, wait for the mic, tell us who you are, and remember the questions end with a question mark.

AUDIENCE QUESTION: Great. Thanks very much. Jeff Alexander. I'm Director for innovation policy at RTI International. And this is for both gentlemen Ray. Thanks again for sharing your thoughts. I'm thinking back to the solo paradox about it showing up and everything, but the productivity statistics.

And part of that is that, you know, its top managers didn't really know what to do with computers for a long time until people figured it out. And then you saw productivity increases. I'm curious what you think is the delta between managerial understanding of AI and managerial, kind of wishing about what AI can do for them.

WESSEL: Right. How long is the gap? Yeah. Anupam in the front here.

AUDIENCE QUESTION: Hi, Anupam. Kana. I used to work at the World Bank. I started off doubling in AI at SRI in the seventies. Wow. But now I'm an economist. Question I have is, you know. There are, I think the problem I have with statistics, and I've been following the AI index ever since it start that and it has evolved, but one of the problems is that there isn't a framework which sort of integrates productivity, the difference, task occupations, et cetera.

I think a lot of progress has been made including by those. But you believe that by now you have the right framework which gives you consistent answers both on the productivity and the labor. Just to give you a point, I've just come an hour ago from a presentation by Chad, not Chad, Ben Jones. Chad Jones is here, and he has basically pointed out some basic diff problems that we have with how to in both these matters.

I can't take, I don't want to take much time. Right.

WESSEL: Is there another one that gentleman in the back on the aisle in the blue shirt.

AUDIENCE QUESTION: Hi I'm Andrew. I'm a co-founder of a startup here in DC that works on AI policy stuff. I'm curious if you have data on whether AI is actually impacting offshoring before it's actually potentially impacting onshore resources.

Hmm. Mentioned software engineers. You know, there's a question of would you rather have someone sitting next to you running stuff on cloud code or someone on the other side of the world and whether we're actually not seeing the impacts here yet 'cause it's actually impacting things in India and the Philippines and elsewhere.

WESSEL: So we don't have time for both of you to answer each question, so pick one. And these are,

RAGHAVAN: So you mentioned, sorry, I'll, I, I'll do really quick. The Ben Jones, by the way, he's a professor of Northwestern. Right, right. Fabulous paper. I just talked about his paper yesterday in a class I was teaching. So I think I think we're still struggling with those frameworks quite honestly, because a lot of the frameworks use AI to evaluate one shot tasks.

And we have to, and this is by the way, it's challenge, it's hard, like how do you evaluate AI in a multi-shot, multi kind of step? Environment is really tricky. And so, you know, I think a lot of the AI labs including, you know, Microsoft and LinkedIn, we're trying to think about that.

WESSEL: So, Ray, do you have any sense of how, how well you're doing it, measuring the the, the framework for measuring productivity and task and all that relative to when you started the index?

PERRAULT: I mean, certainly what we report on is considerably more sophisticated now than it was, but we take no credit for that. That's because, you know, economists have worked on the problem. They've found it valuable and they've you're, they've, they've produced the papers. All we do is point. The readers of the report and

WESSEL: don't tell to yourself short,

PERRAULT: no,

WESSEL: you could be synthesizing all this stuff in a way that someone like me can understand is a very impressive accomplishment.

Nobody could possibly read all those papers and know what, what the point was. Do either of you have the offshoring question is a good one. Is there a possibility that this is showing up more in the offshoring, in in Bangalore, more than in Baltimore?

RAGHAVAN: Well, in terms of hiring performance on LinkedIn hiring rate, I'll just speak to India for example.

That market is doing quite well, even on the software engineering side. And so I would say it's doing better actually than the US market. We have some of that data that's, it's actually published back in January. The data's a little bit dated right now, but we see some growth there. Absolutely.

WESSEL: Great.

Well thank you both very much. Appreciate the, the questions and we'll now turn it back over to Derek.

Appreciate, thank you.

Ask that question. It's very handy answer.

BELLE: Alright. Thank you to our panelists. What a fascinating conversation. We will now move to our last segment of the day which we'll kind of pick up on the threads of data availability, societal impacts, and AI governance and expand the lens to another set of key policy issues around measurement, evaluation, global governance and public opinion as it pertains to AI adoption diffusion.

So the panel will be moderated by Sanjay Patnaik, who is a senior fellow in Brookings Economic Studies Program, and the director of our Center and Regulation Markets. Sanjay, go join us on stage with panelists.

PATNAIK: Everyone. Hello. Here we go. Hi, everyone. It's great to be here and have such a distinguished panel here today with us. We have Yolanda Gil, who is the professor of computer science and spatial sciences, and principal Scientist at USC Information Sciences Institute at the University of Southern California, and she's also the co-chair of the AI Steering Committee and an affiliate fellow at Stanford, HAI.

We have Russell Wald, executive Director of Stanford, HAI, and we have our very own Elham Tabassi who is leading the AIET initiative here at Brookings. I think it's a very exciting time when we think about AI exciting, and for a lot of people also anxiety inducing as it turns out. But I think the panel today, what I really wanna focus on are, are some of the issues that people worry about most, right?

And when you listen to tech folks and then you listen to economists and you listen to like labor representatives for instance, there's a huge gap between what they think is gonna happen and how it will affect society. Right? And I think what is really interesting to me is that the tech folks often assume that technical feasibility immediately means adoption and means economic feasibility.

But those are very different. And diffusion can take a much longer time than actually the technology when it's being invented. And economic feasibility can even take longer because firms will actually need to have the business case to be able to adopt it. So with that, I wanna start with questions for all of you.

Maybe we can start with Yolanda. What do you think are some of the most significant trends that you observe in the report?

GIL: There's a lot and we put a lot of care into these takeaways from each of the chapters, which represents a theme like science or economy or policy. And then the overall highlights of the report.

For me. The last year was the first year that we talked about AI and science, and we highlighted a few of the foundation models in science. So, for example, foundation models for time series data. So they use the same approach as you see in LLMs, but they look at a lot of different time series data and build models of how variables change over time.

And this year we actually have reported on many foundation models across various areas of science and, and medicine as well. And what's remarkable is that even though 90% of. The AI frontier models come from industry. All the multi-model language models, et cetera, 90% of them come from industry. The science models come from mostly collaborations across sectors.

So you see universities, industry and government labs teaming up to create a foundation model for astronomy or water resources, et cetera. So that's very significant. We also see multinational teams. Working on foundation models for science. So that's a very interesting aspect of AI that there's areas in society where we are going to see a lot more knowledge investment and expertise in AI that is devoted to these problems that maybe have less of an industry interest.

Industry also builds models for science, but I think that there's this indication of the cross sector and cross national collaboration. So I'll, I'll highlight that about the report, but there's many other things.

PATNAIK: Great. Thank you Russell. Thank you. Thank you.

WALD: Thank you. And thank you to the Brookings Institution for hosting us.

One thing that I will I just wanna say really quickly, I'm recovering from bronchitis, so I don't want people, if I go into a coughing fit to think that I'm spreading a bunch of germs on you. One thing that I think is really important about the index that I always like to caveat and say, the index is a very data-driven report, but we, and we do have a degree of analysis, but I think this is our opportunity as individuals to give contextual side to this.

And what do I see as the most fascinating thing out of this year's report? I say this not in an effort to exacerbate the concept of a race between the US and China. I say this because the fact that I think there's two ways of doing business in a way and two ways of research and development that's kind of happening.

And it has fundamental implications for policy and the fact that. We are looking at the amount of private investment that went in on the US side compared to the chi Chinese side. Now that's not a fair apples to apples comparison because we're not really measuring the public sector side of this. As Sha noted the data challenges, but you see how fast China has closed the performance gap compared to the United States is significant and has significant implications for how we do research and design on this technology.

One point about this, and my hypothesis as to why that gap has closed so fundamentally or so rapidly is because of the fact that China has embraced an open source community that has allowed for them to build onto their work and build on each other and create this vibrant community and ecosystem that is having a significant impact.

But. We'll, we'll talk about this I'm sure later on, but could we do that in the US right now if we wanted to? I would refer a lot of people to chapter nine, the public opinion chapter here. I don't know if we could to that same extent. We are kind of challenged right now because we have what are the numbers that I think are really interesting to me that I keep saying to myself, 84, 38, 80 4%.

Of Chinese respondents had a much more favorable view towards AI products than the 38% of favorability in the United States. There's a lot of reasons for this and we can discuss this more, but there are gonna be deep challenges related to that. And so I think it's really important to understand these two lanes of development and what that the reverberations and policy implications will be.

PATNAIK: Thank you. And I think this, this last one is especially interesting 'cause the US has been at the forefront of adopting new technologies in the past, but here we really see a deep reluctance. Even the European countries, some of the European countries are more optimistic about AI than the average American, which is very interesting.

Elham.

TABASSI: Yeah. Thank you. I want to underscore, double click on three of the findings that has already been mentioned. One is the decrease in the number of the talents and developer to us. So 89% since 2017 is a remarkable drop. The other one that I wanna mention that Russell just talk about this is the open models.

And I think the report talk about since 2014, we are seeing a opening in the performance of the open and closed models. And of course, open models models that can be better kind of, look inside and test them has a important implications on the diffusion and adoption. And the third one is what Yolanda talk about and it's adoption or actually using AI in science and r and d.

That's really warms my heart because it's one thing that for us to have a very powerful tool on our phones, but it's really, it's a much more important one that we actually get it into advancing. Think of the scientific discoveries. The last thing I want to say is that. This is the report, as all of you know, is a sort of a snapshot in time of the past year summarizing the data and showing it out.

And ai, as ASha mentioned, is a very dynamic field, and it's moving very fast. Also a lot of these trends that we talk about is already being tested with the with what's happening around the world. And example of that is, whatever is happening in Persian Gulf is introducing uncertainty to energy, to infrastructure in a region that we're becoming increasingly important in AI investments.

And I say that as a, as an as a reason for importance of doing this report annually to get this data on the trends, but also what's happening around them, and the backdrop and the context that all of these things happening on. Better understanding all of them, and in a way preparing ourselves for a AI enabled future in a way that we should.

PATNAIK: Great, thank you. I wanna move on a little bit about that clear gap that we oftentimes see between kinda like the tech industry on the West coast and policy makers here in dc. Whenever you hear them speak, it's really interesting, kinda like what they see for the future, how, how, how domestic the tech industry is, how like not concerned many of the policy makers are.

So I'm very curious, like Russell, when you look at some of the information, the report and I think the immigration drop is really, really significant and probably not very well understood here in dc. Do you think this information makes it to the relevant policy makers in dc? What can be done to bridge that gap and bring attention to these really critical issues for our competitiveness?

I mean, this is the economy of the future.

WALD: Yeah. So I've been in this field roughly since 20 17, 20 18. Before it was even called AI policy. I remember that's around the time I met Elham. And she was a, a leader in this field, in government, I think that one of the biggest things is AI literacy and the literacy rate that we're trying to inform policymakers, and they're making the best decisions related to a lot of this, and there's a lot of misinformation that's out there, and it's hence why we are here and have come from California to be part of this and to engage.

It's one of ha's earliest fundamental parts of the organization. Our co-founders, John Etch Mendy and Fefe Lee early recognized that this would be such an, a powerful technology that

it would affect not just governments, but governance itself. And it required to have a policy arm to engage with policymakers.

But indeed I think. That particular stat statistic is very important because that will affect the decisions we make and what can happen in this particular space. But it is and is why the index is such an important part of the conversation because this is data that is very, as Yolanda said, we take a lot of care into this.

We spend a lot of time thinking about this data and we try in an academic space to be as objective as possible and look at this data without the same level of interest that other maybe frontier Labs that have a financial interest or things related to that. So we want as much objective data as possible.

To take it one step further, I think what we need to do is find ways to have this shoring up nonprofit AI research and academic AI research in a way that we can be a good. Part of the system. So it's nothing against industry. I love some of the things that industry is doing and it's great, but I don't, what we need is more balanced ecosystem that's allowing for academia in particular to have a much better robust seat at the table to be as important and influential to policy makers.

Previously before becoming executive director of HAIL was our policy director. And in that role I spent a lot of time combating a lot of misinformation of where a policy elected official will tell me why I read someone's substack. And that is where I was able to learn all of this with no proper citation, no real validation.

And I think that there is indeed concern about the future of ai, but we need to make sure that we're as, my boss. Faith AI often says science, not science fiction. And that's a really important part of the process.

PATNAIK: I think you really pick up on a very, very important point, which is the availability of data and the ability to measure the the impact of these systems on society and to evaluate them.

So when we look in you mentioned it in your, in your speech, right? Like it's quite distressing that the transparency is going down and, and the availability of, of insights into these systems. When we look at what information is out there on the capabilities of these general AI systems and, and what effects they will have in society, we see that good information is hard to come by, especially for researchers, right?

For independent researchers that don't have ties to industry. But it's really important to have that. We need to be able to assess how these systems work what kind of impacts we might see, and then plan accordingly. So, Yolanda, a question for you, right? Despite the challenges that we have, both policy makers and businesses, and even end users, to some degree, we have to make governance and deployment decisions in a rapidly changing environment with varying complete information.

What would it take for relations to produce reliable enough information to effectively inform these decisions? And who should be setting the standards for this?

GIL: So, you know, the report is really focused on data and we, we strive for that, but there's a lot of things that you can read in between the lines.

So, for example, we talk about open models and the importance of having open models. Isn't it curious that even closed models are advancing so fast in parallel. Without communicating to one another their latest discoveries. Right? So there's six companies that report shows that are going head to head advancing AI very quickly in various measures.

So that's very curious that the technology is such that they're all relying on this transformer architecture that came out in 2017 and they're all expanding in different ways and it's, you know, from the point of view of an AI researcher of 40 years, this is miraculous. I mean, I expected a plateau. I've been working with the index for the last two years, but I expected a plateau at some point, and it's not happening.

So this is quite remarkable that it is a technology that the core idea is the same, but the improvements are exponentially better. So you read between the lines of this and there's actually the technology is so powerful that no matter what each company is trying to do to improve it, it does improve.

And it's incredible to see that, that in parallel they are improving this. Now the report doesn't say, but there's very few people in the world, only a few dozen people in the world that can do this, that can advance this technology. And they live in these labs and nowhere else. And you know, it's, it's really, you know, an incredible thing to watch. Those are the workforce and the people that you want to keep where they are. And that's the kind of talent that you want to make sure we have wherever the next big technology shows up. So if the next big advance is, I don't know, quantum machine learning, it's not gonna be that few dozen people that are advancing things on those six hundreds.

It's gonna be a completely different set of people, and I, for one, would like them to come. That for them to come out of our US universities and help us improve the US landscape. So how do we do that? We do that by investing in our you know, academic research environments. We do that by training undergraduates to think in creative and innovative ways, something that's very unique to the us.

So, that's the kind of workforce preparation and the kind of anticipation in policy that I would worry about. I think that you can very easily underestimate how little we know about how to do AI policy. Yeah, because in other areas we've been thinking about policy for decades, right? So I work a lot with safety engineers, people that do medical device safety health safety, patient safety, but also aviation safety, engineering safety in general.

And I see that they have a methodology. They have an entire suite of ways to measure, to consider, to understand a complex system. We have nothing like that in ai. And so we cannot expect that a technology that has had these very incredibly advancing ideas that are not stopping. And then in terms of measurement policy, understanding these complex AI systems, we're really not advancing.

We're really not investing in that. We're really going very. Slowly. And so I cannot say enough about how little we know about how to characterize AI system behaviors, how to characterize these complex emergent learned behaviors, how to be able to measure them, how to be able to characterize when you going to a risk state, a risky state, and how to prevent that and how to manage that.

So, we, we really are asking for a technology that you want to have a policy that makes sense to regulate that technology. And we are really in the very early stages of understanding really how AI works. It, it kind of sprang on us. And so, I just wish that we invested more on the understanding and the, the science of measuring ai.

PATNAIK: And I think an added challenge here is that. Even for regular industries or technologies where we know how to do things right? Regulators operate quite slowly. They take their time. We have the processes in place in the government, and now in the AI systems, we have a technologist that change within like three weeks, right?

And the market moves so fast, so regulators can catch up even less. Elham, I want to get to you because when we look at a lot of these AI systems, they operate increasingly autonomously. So do, do you think we need to reevaluate like evaluation standards and rethink them as kind of like everything becomes more automated without human input?

TABASSI: Yeah. Thank you for that. If I just can add to what Yolanda said about what we should definitely learn from other industries and they learn it over the decades from the incidents and, you know, costly mistake that happened and, and we don't want that to happen to ai. So it's, it's, it's a lot more complicated.

And then we are doing that. We want to do it for a technology that, as, as Sanjay said changes every. Every, every, you know, a few months or so. So, the question about and Yolanda very rightly said that we we don't know enough about to how to measure and characterize this AI systems and the international AI safety report that came early February, also pointed that as evidence gap and talk about how pre-deployment testings are increasingly failing on predicting the functionality, trustworthiness behavior of the models in the production, in the real world context of the use.

And as we are grappling with this agentic AI is coming. And, you know, that's, that's that type of autonomous systems that Sanjay is talking about. And the report talks about the evaluation, landscape benchmarking and the gaps that exist in, I think chapter two and three of the reports.

The, all of the chapters are good roots about, particularly chapter two and three because my fascination with the measurement it talks about and Sha mentioned that in her presentations, that the benchmarks saturate quickly, data contaminations and a lot of other dose problems coming with that.

But I want to say is that agent TKI evaluations, the problem with that is much harder, much more difficult. It's not just about designing the good benchmark and operating a good benchmark, but the question that is our methodology, as Yolanda said, and the whole infrastructure that we are trying to rapidly build for LLMs can even be transferred to, to agent think ai.

So many of us are are definitely me, still very nostalgic to the time, the days and times that we are evaluating technology and the models that were purely static deterministic, and we could measure them completely traceable LLMs challenge there. But I want to say that with the advancements in the evaluation or more or less evaluation of LLM and the LLM themselves are traceable.

You know, more or less you had the one input, one output, one defense answer scores. The models are stateless. Human is in the loop. You can. Rerun the experiments, aggregate results, and analyze them and write a report. Agent DKI is challenging all of these assumptions in many different ways.

The first one is the unit of evaluations is no more just a response. We are dealing with assessing a sequence of decisions tool calls environmental interactions that may unfold over the time. The second is that the environment itself is not fixed. Agents act in environments that can change in response to their actions.

So for the same agents starting at the same identical starting point, they may come to the outcome completely different. So, reproducibility, which was a hallmark of the rigorous scientific evaluations is structurally very hard or almost impossible. The third challenge is that. Failures compound and accountability is, is very hard to kind of, follow or figure out where it is in a multi-agent pipeline.

Agents talk to each other without any visibility into what's happening. So, the if one of these models in the pipeline make a mistake the erroneous answer become the input to the next ones. There is no visible thing into this. So by the time failure surfaces, tracing that requires a full history of the interplay across the agents which is very difficult, if not completely impossible.

I think chapter three talks about responsible AI evaluations and talks about something that in AI or MF we also talked about this, that we cannot optimize simultaneously for all the trustworthiness characteristics. So there are trade off between safety, between, and accuracy between privacy and security.

So these models are being sort of built to make one of them more prominent. So what happens when they are talking across, you know, in a distributed sort of systems, where the where the safety and security is really the weakest is determined by the weakest link in the chain. On top of all of these things is the issue of the situational ever a or simply deception where the agents are behaving differently when they know they are on their evaluations and tests.

So all of this as I think, again, employed in the report is that this is not a benchmark gap. It is that we just don't have the right evaluation paradigm, the right methodology and the field of metrology, which is the field of how to do measurement in a trustworthy way was built for physical measurement.

With all those, you know, traceability and reproducibility elements that are being increasingly challenged for the AI systems nonetheless agree with Yolanda, that we have to study the other field and figure out what they work with them and the things that we have to focus on. On, on, on making them improve them.

PATNAIK: Great, thank you. I wanna move forward to another topic which is getting much more traction, which is AI governance. So we see a big number of players at the international level, at the national level, industry groups, NGOs, getting into AI governance and really starting to think and try to advance thinking in this area.

Brooke, my colleague at Brooks have called, is like a network of networks that can be really good to like generate information. Russell, how has the governance landscape shifted since

the 2024 report and what do you see as positive developments and what is more negative ones?

WALD: I think there has been, just as AI has changed and fundamentally changed our lives, so has the governance related to that and it, there has been one thing that I will take, since the 2024 report is a lot more people are informed and have become much more informed. You can just tell by the degree of questions that are coming at us, even on this kind of road show that we're on, that there are much more sophisticated and nuanced than had been in the past. And that there is a degree of strong interest on, I remember pre-chat, BT Chat, GBT, walking Capitol Hill and there was interest in AI, but there was not much depth to the the question.

So I think people are working to become much more literate in this space. At the same time, I would argue that I think there is still a heavy degree of misinformation out there and not deep rigor in what is behind some of this or what is being discussed in 2023, 2024. We are spending an enormous time talking about compute thresholds and whether something's at 10 to the 24th or 10 to the 26th on compute thresholds.

And here we are in 2026, not really talking about that. So the question is, was that an issue in 2024 or not? And so, I think that we still have a lot of this ambiguity related to this, and it creates a significant challenge. We at HAI have set for the, since 2022, we brought it back.

COVID had created a gap. We initially had it in 2019. Congressional bootcamp for AI. We had an, we have an AI bootcamp for congressional staff. We bring out to Stanford University, a bipartisan delegation of congressional staff. We can only have them for 72 hours because that's all the ethics committees allow for us to have.

But we spend a lot of time teaching them an interdisciplinary perspective related to this. That is one important part, and that I think is important in the report, is this interdisciplinary site to help build out that literacy rate and help people become much more informed on this. We do note, I will say this, we have an entire policy chapter and we note where you're seeing as Sha noted AI frameworks are growing national frameworks are, are coming.

We're seeing more laws passed, more hearings. So there is a national and international dialogue that's happening on this technology. And that is a healthy thing. But as Fefe Lee often says, this is a civilizational technology, and so therefore it requires an, an, an entire part of society, a whole of society to be greatly involved in this.

And that includes policymakers and governments.

PATNAIK: I love that. Last point, I think. And a lot of people don't realize how monumental this time is, what kind of changes we will see, and this is probably unprecedented to what we've seen in the past in human history.

WALD: Yeah. If I could just build on one point, Yolanda's earlier point on the scientific discovery and what is coming from this.

You can't just keep it isolated to artificial intelligence and say, well, this is artificial intelligence. One of my favorite anecdotes here is the nuclear fusion breakthrough that happened a few years ago within the US and it's this, we refer to it as a nuclear fusion breakthrough. That would've not have happened if you did not have AI and ML techniques.

I think you're gonna see a lot of emerging breakthroughs, and we won't call them AI breakthroughs, but they would not have happened if you didn't have AI. And I think that's an important part of the science chapter and it's important part of what Yolanda is trying to put out here related to that.

PATNAIK: That's a really great point. It's a multiplying effect that AI has, right? Elham? How can we ensure that? We integrate a lot of diverse viewpoints and networks in our future AI governance because I think what we want to avoid is that it's being dominated by a few players and then they ignore a lot of the other viewpoints that are out there.

TABASSI: That, that's, that's really important, important questions. And, and as, as, as mentioned, there is a chapter with nice tables that talk about a lot of the activity that's happening around the globe. There are more of the countries having AI strategy, the signaling that this is important and putting thought through this standards, frameworks dialogues in an OECD across.

And that's definitely something to be celebrated and some great thing that's a lot of people are thinking about this also making sure that. These are at least coordinated. They they are aligned and on the results and the outcomes that they come to are not contradictory. So, in a way I think many of the conditions that can ease and facilitate inclusion may also be some of those that can cause if the, if not done right can cause fragmentations but going back to the report, I think Yolanda mentioned that, and Sha showed in her presentation that 90%.

Of the frontier model development are happening in private company. If I said that the numbers right, I think this is this is something to sit with and, and think about this that talks about the information oximetry the open model and open weight. Particularly considering that the open weight momentum is happening outside of the US is also another thing to think about when we think about the diffusion and also.

The more information about the model, more thing that we can look under the hood of the model allows for better understanding advancing the measurement paradigm and what happens with those if those expertise or not in the US. And all of these conversations are happening in the backdrop of conversations about sovereignty and nations are racing towards sovereignty, which is sort of defined as the complete control of the full AI stack.

Which in a report that we put out, I think early February or end of January, argued that it is structurally infeasible for almost any country. And, and instead of going for full sovereignty, thinking about managed interdependence might be a better way of looking at that because AI is trans national and the choke points on mineral chips, energy, data, model talent.

And all this is, is sort of across across the whole point. So, how to reframe the conversations about inclusion. That's that are about mapping the dependencies. Diversifying suppliers embedding interoperable solutions and building those interop solutions through standards of governance and procurement.

So having a seat and a voice in the conversations is not about building the best frontier model, but around the governance and, and diffusion and adoption of all of those things. So, so I think again, it's it's good that these conversations are happening and we should be involved in participate in all of those, but also thinking about how to make sure that we are not exasperating the fragmentation problem.

PATNAIK: Great. Thank you Ivanov into audience q and a, but maybe very briefly, Yolanda and Russell, like 30 seconds each. I wanna talk about the adoption, competition and public opinion piece. What does it mean for kind of adoption and our competition with China, that Americans are much more pessimistic about AI than the Chinese or the Europeans?

In fact, what could we do to change that? Maybe Yolanda, very briefly, and then Russell.

GIL: I think we just have a very open press, a very active and imaginative news space. We have a lot of people online posting a lot of things and a lot of opinions about ai. And you know, as Russell was saying a lot of misinformation.

This is a field that is victim of a lot of misinformation. So, it's also a field that it's been a science for many decades. So if I tell you what year the first Neural network drove a car across the United States, what year you think that was? It was 1997. That's many decades ago. And so there's so much misinformation.

There's also. You know, a lot of questions that we don't have answer for because as I said, we're really behind on the science of understanding ai. I think that it's impossible to predict and to figure out what will work well. So what I think, you know, I'm not a policy expert, but it seems to me that mistakes will be made.

But there's things that we have to act upon and do things on. So, you know, feels like medicine. You know, bioethics is a very advanced field compared to ethics in ai. So we can take a lot of lessons from, from other areas where we've done policy before and it's better to have it than not to have anything.

WALD: 30 seconds is could be, this could be an entire panel. This one little part, I would say enough really quickly, I think, you know. The one thing is, is a, a hypothesis I have here is why this is the case, is indeed what Yolanda is referring to is a very free press. But at the same time, we have a press that's very click bait driven.

And the reality is, is you hear one of two things in the American press about ai, it's either some kind of fairy dust that's gonna fix every possible thing, or it's gonna kill all of us. And because of that, we're not getting the nuanced argument that is desperately needed to have a discussion about this.

And people therefore have this extraordinary fear from, from, from this. And there are so many extraordinary, amazing things that are coming from this technology that are gonna fundamentally transform humanity for the better. And yet there are things we need to be cautious about and we need to be careful about setting rules of the road that are gonna allow for us to be able to get there.

So I, I, I think that's the biggest thing. And I think we need to make sure that we don't have this extreme. Perspectives in our society.

PATNAIK: I think that's a great takeaway and most people probably don't think about it when they get on an airplane that about 80% of the time the airplane is being flown by the autopilot and not by the human pilot.

Alright, questions? Let's see. Gentlemen in the front, please if someone can bring a microphone. Thanks. And please keep it very brief. Just introduce yourself and briefly ask the question.

AUDIENCE QUESTION: Hi. Michael Dilla from Americans for Responsible Innovation. I want to ask a quick question about narratives, about races.

The China race, on the one hand, the competitive race between American AI companies on the other, you each Yolanda and Russell made comments about what data show and don't show. So I'm curious what you would say that the data tell us about the value of those narratives, because it seems to me that on the point of the exaggerations of the press, that both of those narratives are rather durable within.

Thinking communities, and I'm not sure that they're very valuable. But curious what you think the data tell us.

GIL: I'll let you take this.

WALD: I'm not sure I completely, I I, I'm tracing the question in terms of the fact that I do think, the narratives that I see in terms of a race, and I wanted, when I started my point off is I don't want to intend to add to a race and, and have that, but competition can be healthy and can be a good thing. We can learn from others in a competitive environment. What to do that concern is, is what are we in a race to the bottom and what does that mean?

And I see in many respects that that is the case in both the geopolitical context that you're referring to it as, as well as the internal model development within the domestically within the United States. How do you avoid a, a race to the bottom I think will be confidence building measures at the geopolitical level of where we can actually start to have a alignment between that and where we can have the appropriate and right regulation that allows for the innovation to happen.

And not I'm, because I know there's a time, I'm happy to talk to you more offline when you're referring to the data of that. We could always use more data to inform us, but I I, I would need more. Context.

PATNAIK: Thank you. Another question over there please.

AUDIENCE QUESTION: Thank you. My name is Shabnam Moti and I have a question about the data and also the digital literacy of agencies in, in DC specifically. So a lot of the hype around AI has led to quick adoption including after the Doge cuts within. Federal agencies. And so, I I'm also wondering about the loss that that can't be measured easily, the loss of economic value due to AI being adopted too quickly.

We also see lawyers being disbarred because of that very reason. And and so I wonder if you're also, Russell you talked about working with with Congress, but are, are you also doing work with any of the executive agencies to to kind of. Have a more measured approach to this adoption?

WALD: Yeah. We absolutely work with agencies. There's we don't have a formal, we actually have had a formal training pro program for agencies. I should correct myself. We

partnered with GSA and HAI is trained over 8,000 civil service employees through a spec, the AI Training Act, of which we did the first lead on in the first year, and we've done a second year of it.

So indeed, we have been trying really hard to train agencies and to be part of the conversation. And it's not just in terms of trainings it's also briefings. We've come from briefings this morning at OO. Other agencies Ray and Yolanda have done other briefings as well since we've been here. I think this is in my view, a.

Social responsibility we have as universities to go beyond the AC academic campus and to be able to inform people as much as possible. I think it's why you have think tanks like Brookings to be able to inform people, but that is really important, that objective analysis versus industry driven analysis and where agencies and policy makers need to be able to discern the difference between the two.

GIL: I wanna say that AI and, and almost any technology that you introduce in an organization there's well known methodologies and principles to do that, and they're just not followed everywhere. And so you may be seeing the effect of that.

PATNAIK: The gentleman in the back with a purple shirt, please.

AUDIENCE QUESTION: If I may politely re-ask the question on fear of adoption or fear of adoption justice Shannon, NCIA, sorry.

What is the role with engagement with Americans in AI impacted communities to get to know their concerns, specifically environmental data centers, adoption threats of employment threats of unemployment or AI chatbots in order to try to overcome this concern of Americans being more fearful than peer nations in an environment where there are people who have specific niche knowledge bases.

You have people who are afraid of data centers in their community and have very specific knowledge on that. You have people who have very high environmental knowledge that are pushing back against AI. What is, is there a way to, is there a role for AI companies to address this? Who else could help address this and what should be done to address.

PATNAIK: Anyone wants a take?

GIL: I'll say. So, so I think this technology spans so many aspects. So first of all, these are companies that have our data and so we have mixed feelings about them. They also make a lot of money that we don't get, and it's, you know, because of us and we don't get any of it.

So we have mixed feelings about that for sure. And then we hear about them consuming. Energy affecting the environment, et cetera in ways that, you know, really bother us. But ask the individuals that use AI every day on their smartphones to curtail their use of AI so that water can be saved, or energy can be saved, or that these companies don't have their personal data.

And none of us would do that. And so I think that as individuals, we have this love hate relationship with technology companies and AI companies that we really need to digest and

come to terms with. So I think the public opinion aspect and the the community aspect is very complex space. What I will say is that.

Not just AI and, and I'm an AI researcher and a computer scientist. Computer scientists have done a horrible job of helping the public understand computing. So, you know, you listen to NPR Science Friday and Ira Plato is an excellent, excellent interviewer when he interviews a computer scientist. He's humming and humming.

He can be talking about mushrooms and fungi and quarks and anything. But when he comes to computer science, he's just so uncomfortable you can, you know, see it. Right? And so this is a super smart guy that should be able to talk about technology without reimbursement. Computer science. No, you go to Wikipedia.

The computer science pages is like the first sentence. You cannot follow what it's saying already, right? So a technology that is. So fundamental to our lives, you know, as, as computer scientists, we've done a very poor job of educating the public, conveying what it is, how it works, what it, you know, all of that.

So with AI, it's the same story all over again. We're just not helping understand things. I think one of the most fundamental things that I, I tell when I talk to public spaces is that. My hope for AI is that it will help us become better people because it is exposing very clearly and quantitatively our biases and the ways that in which we are dealing with our personal spaces, our personal data, our personal questions, our personal tasks you know, when students cheat, when, when there's all kinds of irresponsible use of AI out there.

So I think, you know, my hope is that it will help us, you know, turn the mirror back to us and help us. Become better people. So that's a message I would say. And don't be scared if Wikipedia and other sources don't help you learn about AI. Make a point to learn about it. Find some way to learn about it.

There, there are many ways to do that and people need to know how, how this is. And just

PATNAIK: to add on your question, especially in the data centers, we're doing some work here at Brookings. I think it goes back to your Misco misinformation point earlier. You need transparency in communication, right? And you also need to convey to the public what is the economic counterfactual.

We all want economic growth. We all want to power our engine of, of economics that we have had in the United States for many decades. So you have a choice, right? Like, would you rather have a data center in your community, get tax revenue through it, et cetera? Or would you rather have a chemical factory?

Just think about that, right? And I think there needs to be like a weighing of different options and alternatives. 'cause not having growth is not one right? Alright, with this, we are at the end here. Thank you so much to the excellent panelists. Thank you to Derek. Over to you.

BELLE: Yeah, absolutely. Alright everyone thank you all so much. Thank you to, to Sanjay, to our incredible panelists. And thank you to everybody who joined us today, both in person and online. Thank you also to our incredible teams here at Brookings and at Stanford HAI who made all this possible. Today has been an excellent discussion.

I learned a ton and hope you did as well. I wish I had something profound and interesting to say, but I just say, please read the report, engage with the data. Think about ways. That, that we can all make sort of more empirically data-driven policy possible. And any of the projects that Elham and Sanjay have been talking about, you can find in [brookings.edu](https://www.brookings.edu).

Great website. You can find the report at the link in Sha's presentation. And for everyone who is still here there is refreshments outside. We brought entirely too much. So please go drink coffee and eat pastries.

Okay.