

FINISHED FILE

BOSTON UNIVERSITY
AI, PUBLIC HEALTH, AND ETHICS
OCTOBER 14, 2025
1:00 P.M. ET

Services provided by:
Caption First, Inc.
P.O. Box 3066
Monument, CO 80132
+1-719-941-9557
www.captionfirst.com

This text, document, or file is based on live transcription. Communication Access Realtime Translation (CART), captioning, and/or live transcription are provided in order to facilitate communication accessibility and may not be a totally verbatim record of the proceedings. This text, document, or file is not to be distributed or used in any way that may violate copyright law.

>> DEBBIE CHENG: Good afternoon. My name is Debbie Cheng, and I currently serve as Assistant Dean of Data Science and the founding Executive Director of the Center for Health Data Science at the Boston University School of Public Health. On behalf of our school, welcome to this year's Bicknell Lecture.

Thank you to the many who helped make today's conversation possible, including the BUSPH Dean's Office and Communications Team. Today's lecture is co-hosted by the Boston University School of Public Health and the Center for Health Data Science.

Each year, the Bicknell Lecture provides a vital forum for thoughtful dialogue on the evolving challenges and opportunities in public health. Dr. William Bicknell envisioned this gathering as a space to share and discuss bold ideas—where iconoclasts and original thinkers challenge convention, spark reflection, and inspire renewed energy and commitment.

Today, we honor Dr. Bicknell's legacy through this conversation, and we are especially delighted to welcome his wife, Dr. Jane Hale, whose presence makes this occasion all the more meaningful.

In that spirit, we turn our attention to one of the most urgent and complex issues of our time: the ethics of artificial intelligence in public health. Today, our speakers will explore

how AI is transforming health systems and decision-making, and will discuss the ethical frameworks needed to ensure its responsible and equitable use. We are honored to have you with us for what promises to be a thought-provoking and important conversation.

I now have the privilege of introducing today's moderator, Krishna Udayakumar. Dr. Udayakumar is the founding Director of the Duke Global Health Innovation Center, where he leads efforts to scale and adapt health innovations and policy reforms worldwide. He also serves as Executive Director of the non-profit organization, Innovations in Healthcare. At Duke University, he is an Associate Professor of Global Health and Medicine and a Core Faculty Member of the Duke-Margolis Institute for Health Policy. Dr. Udayakumar chairs Duke's Global Priorities Committee and has published in top journals like the New England Journal of Medicine and Health Affairs.

And with that, Dr. Udayakumar, I'll turn it over to you.

>> KRISHNA UDAYAKUMAR: Thank you, Dr. Cheng, for that introduction and thank you again to everyone for joining us today.

I'm the moderator for this session, we have speakers who will have 8-10 minutes to provide opening remarks and we will have robust discussions about what is driving the future of public health but also economy and society, more broadly.

I will introduce all three briefly up front and we will go into their presentations. Please feel free to send questions using the Q&A opportunity through Zoom. And we will make sure we incorporate questions from the audience as we go.

We know this is an incredibly important topic and one that many people are interested in, so we will try to make sure we keep the conversation moving and incorporate as many of the questions we get as we can.

First we are going to hear from Nicol Turner Lee. Dr. Turner Lee is a Senior Fellow at the Brookings Institution and Director of its Center for Technology Innovation. She is a leading voice on equitable tech access and AI bias and founder of the AI Equity Lab. She frequently advises policymakers and global institutions and her work bridges research, advocacy and public policy to ensure technology serves all communities.

Second we will hear from Suresh Venkatasubramanian. Dr. Venkatasubramanian directs the Center for Technological Responsibility, Reimagination, and Redesign with the Data Science Institute at Brown University, and is a Professor of Computer Science and Data Science. Dr. Venkatasubramanian's background is as a computer scientist and his current research interests lie in algorithmic fairness, and more generally the impact of automated decision-making systems in society.

Finally, we will hear from Kay Firth-Butterfield. Ms. Firth-Butterfield is CEO of Good Tech Advisory. As a global leader in responsible AI, she was the world's first Chief AI Ethics Officer and formerly led AI and Quantum initiatives at the World Economic Forum. A TIME 100 Impact Awardee and Forbes '50 Over 50' honoree; Kay advises international organizations on AI governance and ethics.

We have a wonderful conversation coming up and I would like to invite Nicol Turner Lee to get us started, please.

>> NICOL TURNER LEE: Here I am. Well, first and foremost thank you, Suresh, for that invitation and for those of you who are part of this lecture series, I'm humbled to present before you.

I would also like to say that the AI Equity Lab, for people wondering what that is, we are a lab focused on interdisciplinary cross sector civil society engagement in ways which we are creating purposeful, pragmatic and nondiscriminatory AI.

I can talk more about it in the Q&A but that's a pet project I've had for some time. I will also share some slides but before I do, I want to suggest that I wrote this beautiful book as well. "Digitally Invisible".

In my presentation I will reference a story, because there are stories about people not connected to the internet but one that is relevant to this in healthcare.

I will share my screen, if you all don't mind. I want to make sure as we go into the conversation that we do some level setting because you will hear from my esteemed colleagues. All of us which probably have different perspectives on the same topic.

So let me get started. First and foremost I want to make sure before I start my remarks to do some level setting that we are talking about an ecology of AI. Obviously, there's artificial intelligence, the big bubble, not only talks about algorithms but also autonomous vehicles and other autonomous mechanisms, devices and software applications.

Machine learning is essentially the computational engine behind AI. There's deep learning which permits much more invasive, I guess, technology scraping when it comes to speech, facial recognition, image, emotion.

A buzzword most of us have been hearing and talking about is Generative AI, which relies upon mostly considered the eyes on the back of your head, publicly curated internet content actually feeding into large language models which permit predictive, summative and now agentic, more autonomous outputs.

So I want to frame that because I think it's really important for the purpose of my talk to talk about AI in

healthcare so you understand how these systems work.

In particular, my research looks at bias. And bias mitigation. And I want to continue, now that you have a general level setting of AI to talk about how I frame the introduction of bias, algorithms in any AI system. It could be introduced when it comes to the design of the program. Whoever sits at the table matters. If there are no health practitioners or very few patient advocate or few in pharmaceutical, the AI model suffers because it defines and designs those products and services in computations and ways that may not be inclusive of impacted groups.

Bias also shows up in the training data. The data that is actually teaching the machine on how to learn the predictive or summative or generative experiences of impacted groups or populations, that are the subject of the AI. There are trade-offs that we make. One, because much of the data that comes from AI systems comes from the internet or datasets curated for commercial entities and they come with human prejudices. Over or under representation. I will talk in a moment how this is so important in healthcare. I call this training data that often comes traumatized. It's based on the historical and systemic inequalities that show up.

In healthcare, think about it for a moment, we know from the representation in training data that Black and Latino patients for example are underrepresented in clinical trials. Essentially that training data represents those marks of inequality that impact that last bucket that show the results. Meaning how groups are affected by AI often tends to happen based on the representation in the training infrastructure.

I want to go back to my book for the sake of time, come to this example. There was a woman I met in a public library, talking about library access and how we give broadband to people who do not have computers at home. She was basically struggling to get a library card. She was in tears with the librarian. I followed her out to ask why, only to find out she had stage IV breast cancer and she needed a library card because it was the only way to look at her doctor's records.

Think about if, a person needing a library card to look up private patient data. Somebody asked me on a panel if this young woman, who again was a mother of three, her name was Frances, had AI help her through the course of her diagnosis where we had discovered it much earlier in its maturation. I had to tell them probably not, because people like Frances, and people like me, Black women are underrepresented when it comes to breast cancer diagnoses. We don't participate as often in clinical trials.

So, I share that with you all because we have this ecology that is still baked with bias. It's important that we understand

that could have adverse effects on various populations, which in the application of AI in healthcare means a lot. It is a fine line between whether or not a person like Frances will stay alive, or will live a life where she is not able to have the quality of care that is needed.

This is important because over the last couple years, we have seen AI in healthcare show up in a variety of areas. I just listed three for the sake of time.

Clinical trials, and of course, let me not be total Debbie downer with this.

There are some parts of healthcare where clinical trials were represented. I was on a panel Friday with the Association of Black Cardiologists, where one of the esteemed doctors said, "Cardiology imaging and screening - we just do a better job when it comes to AI. We have beta tests we are able to look at, based on the fact more people are engaged in this kind of interventionist activity."

But there are cons, based on who is harvesting that data. Where the data is being taken, there are community clinics with less sophisticated radiology equipment, where the image may be more grainy. Friends, it shows up in healthcare clinical trials and algorithms in datasets where we are not accounting for the representation of various groups.

One area is health equity research. While AI in healthcare can help us to get to granularity when it comes to social determinants of health. For example, being able to understand why there's a propensity of one disease over another in a population. It may not take into consideration external attributes.

One of my colleagues wrote a nice paper during the pandemic that suggested that more people of color were dying not only because they were medically underserved but they lived in dense housing conditions. Or they lived in areas that were deserts when it came to hospitalization or quick care and use of the emergency room and distrust of the medical system actually limited them.

If you aren't understanding it by now as I get ready to close in just two minutes, what I'm suggesting is there are breakthroughs in AI when we think of how it will influence the healthcare space. Much of that will go down to whether or not we continue to interrogate AI for what it's worth. Is it demonstrating effectiveness? Does it have situational consciousness? Does it work like I said in radiology, the same way in gastroenterology. Can we ensure the second bullet, it can surpass any trade-offs simply because we don't have enough data or research in that space. And who do we hold accountable when things go wrong?

I'm always reminded when I think of healthcare, a young woman named Henrietta. She died sick and poor because no one was accountable for that type of cell harvesting that contributed to the greater good. In the end, what I would suggest the four things as I wrap up my comments. I'm very passionate, I could probably go over my time but I'm timing myself as I'm speaking.

First and foremost in the healthcare space as we think of AI, let's always consider who is digitally invisible, remember the book I talked with? It applies here too. And who on your team represent the interest or lived experience, whether expertise, in subject matter expertise, a community not represented, make sure that's part of the conversation. Always examine the quality of inputs. We know AI will be a big it factor as we see less people going into public helm, shortages in nursing and doctors and we may need to rely on AI to help us get through those shortages. But they are done with trade-offs on shiny objects and products and services that still create chaos if they aren't properly trained or having transparency with the communities they are impacting.

Build in accountability. And always embed any use of AI particularly in healthcare with responsible privacy, fairness and equity.

I'm going to close here because I think this is the last important thing to say. We do not want to use AI in ways in which we widen the gap of health disparities. So it's incumbent on all of us to be mindful of how these impacts really affect every day patients who may not understand how this technology could actually help to extend their well-being, or physicians who may not understand that in the efficiency of being able to use an AI tool to transcribe notes, that we may be missing something. There was a recent article about that in The New York Times. I look forward to the panel and conversation. But again, let's have a conversation as part of the Bicknell lecture as they have infiltrated the sector.

>> Dr. UDAYAKUMAR: Thank you so much, Nicol, for grounding us in the right terminology. Revealing the sources of bias and ways we may start to combat it. I'm sure there will be much more to talk about in the pieces that you have already laid out for us. To help build on that, let me welcome Suresh to this conversation and hear about the work that he is doing.

>> Suresh: Thank you very much. Thank you to the center for having me. I'm delighted to be here, I also have a short presentation that I will share right now.

As we move from the systems that Nicol talked about to these Generative AI systems. For many years we have had these systems in play. We heard many issues associated with these systems. We heard about the ways in which AI can introduce

racial bias when managing health. We have seen algorithms being used to deny access to care because the lack of transparency. We have seen technological challenges using computer vision to detect for example melanomas.

We have had people, in fact some colleagues of mine at Brown working on better pulse oximeters.

We have a fairly decent idea of how to evaluate and mitigate the problems associated with AI deployments. I think of this as a triangle. We found ways to frame the questions we are concerned about in terms of health equity and disparities and transparency and accountability. We found ways to measure and assess the degree which systems were deploying have problems.

We have come up with tools and interventions that allow us to mitigate, not completely remove but remove the worst effects and therefore get the benefits.

I think of this as having sensors, that's a picture of a nuclear control panel with a bunch of sensors that tell us where things are going wrong. Think of your car dashboard that lights up when things are a problem. We have ways to fix a problem, if your gas indicator goes off, you have to go to the gas station and fill up gas. We understand how to do a lot of that with the systems in deploy in medical field and so many other sectors. Of course with ChatGPT and Generative AI things have changed. We seem to have moved away from a Swiss Army knife where we have specific tools for specific problems that we want to build interventions for to mitigate problems with, to a single sledgehammer that purports to do everything we want. ChatGPT or Gemini or Perplexity or Llama or Qwen or Deep Seek. Or Claude. We just want one tool. We seem to be given one tool to do everything.

The problem is, this is a problem that is well known within the community of scholars studying the systems. Is that once you do that, we don't know how to evaluate them.

I mean, I joke a little bit but a lot of the evaluation of Generative AI systems, the kinds we are seeing deployed right now are what the kids would say just vibing. It seems to work well, we cross our fingers and hope for the best. We don't have any good understanding how these systems work in the field where they are being deploy. Which is kind of scary when you think of the rapid deployment of Generative AI systems. There's a bunch of papers now talking about this. Sort of pointing out over and over and over again that evaluation of AI systems, modern Generative AI systems being used as sledgehammers including public health, we lack the understanding where they work and don't work.

This is a problem, the current administration has a key part of its plan is deploying AI in a variety of sectors

including healthcare. That identifies distrust or lack of understanding is holding us back in terms of deploying it. We want to deploy systems that we can trust that behave in ways that are appropriate and meet our expectations in the way we have had to work with systems in the past.

So how do we get back to what we need? This idea we can frame the ideas we care about and build tools and interventions. The thing I want to talk about a little bit is a thing I feel very strongly about.

That is becoming more and more of an issue with Gen AI systems, that we need to empower the people that work in the domains of Generative AI to do the work of evaluating and framing. Part of this is work we are doing at Brown. Based on a talk, we were asked for hot takes. This is one of them. That basically, what we want to talk about is, how do we bring communities and stakeholders into the process of evaluating AI systems in a way we used to do and no longer seem to be doing with a sledgehammer approach. What does it mean a community-driven LLM evaluations? We want contextualized and specific evaluations. Not just researchers or developers or people at the big AI firms. But people in the spaces where these tools are going to be used, right? Because they have the domain expertise. And would take the lead in decision making. This is really important. Because the people in the domains know how to evaluate systems for specific concerns they have. They know what problems they care about. They know how to frame them, measure them and know what mitigation looks like. Developers and even researchers don't.

This is the kind of thing we don't do enough of and do less of now and need to do much more of.

There's been a number of examples, we have seen this happen for example, evaluating the quality of chat bots for providing information about elections. We have seen this to evaluate the quality of large language models to get reliable legal information. We have seen this being done to evaluate whether these chat boxes are good at scientific papers for journalists and so on.

The fact the studies have been done with community partners and people who have expertise in the area has revealed exactly where they might work well. There's been a lot of calls for doing this. Why AI evals need to reflect the real world. This is something we are moving towards and important to focus our efforts and thinking about how we can do this in different settings.

I want to say why this is important, not just for all the reasons I mentioned. To make sure we are bringing people in impacted by the systems and actually have the judgment to

evaluate the systems well. It's because it actually yields good science. It strips away a lot of marketing and hype we are confronted with AI. It reveals not just binary AI good or bad which I don't think is helpful. But where using AI tools might add value and where they might not. To get a more nuanced question going. Allows meaningful interaction between stakeholders *. not just trust us, we are the technologists and we will tell you what you need. It improves evaluation processes and shifts the power to those who are using them.

This is where are able to properly answer questions like who is getting the benefits, who is not getting the benefits, where the disparities are coming from and so on.

There have been calls for this over and over again. Colleagues of mine at Cornell have been looking at this, I have looked at this as well in other contexts. This is something beginning to be much more of a concern and issue for how we do AI deployments.

And we cannot go forward with proper trusted deployments unless we can find a way to bring communities of practice of expertise into the very design and evaluation process.

So that's what I want to leave you with as a thought. How do we go back to the sledgehammer to a Swiss Army knife that was more effective for measuring and mitigating to get the benefits through community-driven evaluation. Thank you for your time and I look forward to the discussion.

>> KRISHNA UDAYAKUMAR: Thank you, Suresh. What are the right interventions and making sure it's not a one-size-fits-all. Thank you for introduce the importance of community and community-based approaches and also this idea how we think about evaluation and the questions we ask being really important to make sure we are driving the field forward and in an effective way. Again, much more to unpack as we get into the conversation shortly. Thanks so much for setting that up. For our final speaker, let me welcome Kay Firth-Butterfield. The floor is yours.

>> KAY FIRTH-BUTTERFIELD: Thank you very much. It's an absolute pleasure to be here with you today. I'm following very much Nicol and Suresh and I'm very pleased we had this opportunity of them going first and me following up. I'm going to share my PowerPoint. I hope.

Yes, there we are. So what I want to do today is sort of follow-up on really widen the aperture a little bit but focus on AI and health.

I spend my entire life talking to companies, including healthcare companies and hospitals about how to use AI wisely, or responsibly in their businesses.

And as Suresh said, there's so much hype out there. So my

invitation to you today is what do we think deeply about the role of AI, the human spirit and the living planet.

I think, as healthcare professionals that's very much within our remit and it's what makes us human to human so important.

I want to invite you to develop clear-eyed understanding of the responsibility that we bear when we use AI.

I want us to open up a space for each of us to wrestle with this profound question of how intelligence, human and AI will shape every stage of our forthcoming lives and those who come after us.

And I want us to join a conversation together. I think it's so important that we co-author the future that we want to hand to our families, our societies, and the earth itself.

So at the moment, it just feels as if we are running so fast to give away our humanity. All the jobs that we do, oh AI could do that. I'm a former judge. And everybody says oh, AI could do that.

But actually, my response is, well, maybe. But at least, and yes I have biases, obviously I do, I'm a human being. But at least I can explain my judgment. Whereas, Suresh pointed out, if it's an AI judge it can't explain its judgment because we have lack of transparency. I believe we should not be so eager to divest to AI of the things that make us human.

And there are some studies now beginning to come out that show that actually, we are not necessarily improving our intellect or improving our society when we use AI.

But first, I want us to think about, well, what do our patients think about AI? When we are using artificial intelligence, what are they thinking?

Well, 95% of people say it's a bigger number than last year, which is only 57% of people heard at least a bit about AI. They obviously hadn't heard anything of what Nicol told us.

50% of those said they were more concerned than excited about AI.

57% rated the societal risk as high. And only 25% the benefits as high.

53% of Americans are not confident that they can tell a deep fake, nor can I. And I bet all the professionals that are with us today can't either.

And 60% said they would like more control over how AI is used and they don't want it to govern our lives.

So, when you are talking to patients and you are using the next big AI thing, this is the sort of information that you need to have in your back pocket.

Some people you will be talking to who don't know really anything about AI.

And there is some research coming out now. Obviously particularly with Generative AI, it's so young that we don't have any good research.

But we do know that if you write things down, instead of typing it, you remember it better.

There is some research that says undergraduates who use only AI to help them do their research and write their papers actually are illiterate at the end of their degree, their undergraduate degree in the subject matter itself. So that's of huge importance for us.

There's a study from MIT that says if you are engaging with AI you aren't engaging as if you do your own research. And this brand new paper from HRB which shows that businesses with either lazy or over worked people using AI end up with beautiful presentations, or beautiful papers, but somebody along the chain of receiving that within the company has to spend up to one hour, 57 minutes actually making it useful.

So those are some of the big problems that we are already seeing beyond the big problems of bias and transparency.

And somebody earlier actually asked about the planet. And I can tell you that Eric Schmidt said by 2030, AI will use up to 95% of the current global energy. Every time you ask a GPT, one of those Generative AI models you are using a quarter liter of water, so I tend to ask it important questions.

And we are already seeing it overtaking the aviation industry for global greenhouse gases.

I haven't got much time so I'm going to move on quickly to hallucinations.

Hallucinations are where the AI simply makes stuff up.

It's significant in law at the moment. We used to think it was only in the open models like ChatGPT, Claude, et cetera. But we are now finding that even in the proprietary models that have been created specifically from lawyers, we still have hallucinations as to the titles of the cases, the scope of the cases, the facts of the cases, and the decision of the case.

That's hugely important for us to think about when we are thinking about research. When you are thinking about well, how do I use this in medicine, even if you have a proprietary tool. Just recently Deloitte delivered a paper to the Australian government and the Australian government found hallucinations in that Deloitte paper.

What happens then is that Deloitte paper is in Deloitte's proprietary data collection. And so unless it's weeded out, then an AI using that data will be able to find it and repeat the mistake. And I don't pick on Deloitte, it could have been anybody. And it is almost certainly everybody.

And that we call cannibalism.

Deep fakes and the rule of law. My lawyer boyfriend is a case that I heard recently from a colleague of mine. He said, oh, his client said to him, oh my boyfriend says that you're wrong, are giving me the wrong advice.

Well, he asked, who is your boyfriend. And the boyfriend was ChatGPT.

The person had put all of the court cases, including the other side's confidential documents into an OpenAI ChatGPT. So there are consequences of this are massive. And if you then think about your patient putting all this stuff into an Open model, all of their medical records, et cetera, you can begin to see the problems.

In law we are also seeing deep-faked evidence, for example, medical reports in personal injury cases.

And I think this comment by the judge is so absolutely pertinent for our time.

And of course, we are told that everything will be better because of AI. But I think we need to stand back. Yes, maybe the world's economy. But we have not yet as humans, I think defined what is better. Is it how much money we've got? Is it that the economy is abundant rather than scarce? Is it human well-being? Is it a healthy planet, or what?

And the other thing that I want to very quickly, I know I will be out of time and talk about is that the CEO of Anthropic says entry level white-collar jobs will be overtaken by AI. In light of what I just said, maybe we need to rethink this until we get to a point where actually AI could do the job as well as we can.

Which brings me to my question.

Are there any human-only jobs?

And as somebody who suffered from breast cancer, and enjoyed fabulous human doctors to help me with that journey, I would say oncology is one.

A doctor surgeon wrote in The New York Times last year that actually anybody, including an AI, could tell somebody they are dying. I personally think that's a human-only job.

So finally, there's so much to question. I too have a book, but it doesn't come out until January.

But I really hope that this is a good guide from the three of us into the sort of things that you need to be thinking about and I look forward to our panel. Thank you.

>> KRISHNA UDAYAKUMAR: Wonderful. Thank you so much, Kay. Also for reminding us that humanity is really at the central part of this, including how we interact and use AI. Thank you for introducing also really important point around sustainability, as the AI economy gears up. Making sure that we are really thinking about accuracy and hallucinations as an

important aspect of how and when and where we deploy AI.

Of course, remaining grounded on what's better, what is a problem we are trying to solve through this technology? So if I could ask Suresh and Nicol also to come back on camera. We will dive right into the conversation part of this.

So thanks to all three of you for such fantastic talks to get us started. As you have seen from the comments in the Q&A, there's no shortage of really interesting points and where we could go into.

Let me start us with stepping back and really thinking about what this lecture is about and it's around the field of public health.

We are in the midst of profound changes to the field. That includes especially the erosion of trust, of public health, trust in science, institutions driven by so many different factors. And in that, of course, we are seeing the rise, as you have all talked about, of AI.

And I would love to start chipping away at this intersection between the role of AI and the way of impact this trust in public health in positive and negative ways. Nicol, let me bring you into this conversation first.

In New York, how do you think about the ways that AI can help us build trust in health and healthcare. Especially among medically marginalized populations?

>> NICOL TURNER LEE: So this is such an interesting question. And I will be sure to share with people, or you can go to my Brookings web page or my personal web page. I wrote a paper whether or not AI could substitute the public health workforce. And the extent to which with some of the challenges we are having with this particular job that has high burnout, high stress. Could you employ AI to assist with augmentation or complement of their roles, et cetera.

And where I landed, which is where I land with many groups medically marginalized, for example, is to, I think, Kay's point about going through cancer, and having the caring human approach as part this.

So AI is not necessarily going to build the type of trust that we already know is essentially not there with certain groups.

And so I think we need to move that narrative out of the way. Because it is sort of misinformation if we think it's going to positively impact our relationship with patients or people who choose not to seek care.

What we found with the paper that I wrote is that it can do other things like professional development. Or it could help with scheduling, patient management and productivity. But when it comes to that one-on-one relationship, there is a process for

building that trust. And we just put out with one of our non-resident Fellows with the center of technology and innovation a podcast and we are about to release a paper. With this particular Fellow researched how Black women interact with AI. What I found fascinating wasn't that AI couldn't assist with health and employment and other needs but their lack of trust in those institutions really deterred their faith in the technology to work.

So I will just end here. I think it's really important for people to understand that AI is not going to replace, you know, the entanglement of how people interact with institutions already before them. If anything, it could encourage it in some way if the AI works well. In other ways it could discourage or kind of feed into what their current relationship is with those institutions.

So we have to remember that we need a healthy ecosystem already before you start inserting technology into the mix. And I think that was the premise of my colleagues. And it must be very participatory because many products and services being developed as well are not being developed by people with lived experiences of this group.

So I will stop there, but I will share that research we have been looking at, some fascinating space as to whether or not many of these caring economies can actually use AI to build something that has been systemically a long time in the works in terms of trust and relationship.

>> KRISHNA UDAYAKUMAR: Yeah, thank you. Again, an important reminder that any technology certainly including AI isn't going to reverse other underlying trends. And many times reinforcing of those trends and helping them go to scale much more in whichever direction those trend lines go.

Kay, let me ask you, how do you think about this issue of governance and the way it ties into trust, not just in AI as a technology, but the way it is applied and ultimately the field it is supporting, in this case public health?

>> KAY FIRTH-BUTTERFIELD: Yeah, absolutely. As the world's first Chief AI Ethics Officer and spending all my days helping companies think about governance, I will say this, wouldn't I? But governance is the way to trust. If you offer, if the people who you are working with know that you have thought through all these different challenges then they are much more likely to trust you. So I actually now find myself to businesses, well, it's just a wise thing for your customers, your own employees to know that you are governing well.

And when you think about that Pew Research, and there's a lot of research out there that shows so many people really distrust this, then if you are going to use it, you have to

build the trust in some way. So you have to be able to offer your patients that level of understanding of how you are going to use the AI, and why you are going to use it.

I went to the doctor recently and the doctor said, I'm going to use AI to take notes, is that all right?

That was it. They didn't explain what AI was. I might have not known what AI was. They didn't explain anything about how it took the notes or what it did with the notes or anything. And of course, I, being me said, well, I don't mind, as long as you make sure that you check the notes. Because if the notes are wrong, that makes the whole conversation that we are having about my health wrong.

And in an ongoing fashion. I think there are some things we need to think about. As I say, one is public adoption and public trust.

And the other day I was in the toilet of the lady's toilet in an airport and I could hear this woman in the other cubicle shouting down the phone, I just want to speak to a human!

And I think that, you know, in fact United Airlines has now said you can speak to a human. And they are branding. So I think what we have to do is really make sure that we understand that we might be here, where we think about AI. But our patients are here. So we have got to bridge that gap.

You don't want people bringing in their own devices and asking medical questions of an OpenAI that might have the wrong information.

And you know, obviously testing, testing, testing. Red teaming, all of these sort of things. But I want to just finish with something that Nicol is saying and I'm saying. You know, the public don't trust this. But actually some of the public do and are turning to it. So just three stats.

One, the largest group using ChatGPT are white men in America under 46. There is research that shows that 20% of white men living in America are using AI as an intimate partner.

And there was just something published yesterday. I can't remember where it was that said one in five school children had AI as a partner as well.

So generationally, maybe the trust in AI will change. But at the moment, we are definitely not there and we should not be there and we should be really frightened about the statistics.

>> NICOL TURNER LEE: Hopefully, Krishna, we get back to that, because I think it has an application in the medical field.

>> KRISHNA UDAYAKUMAR: Absolutely. Let's do that. Suresh, let me bring you into this conversation as well. You talked a lot about the approach to evaluation and that community-based piece of this. How do you see that really aligning to the issue

of trust and what comes out of some of these tools, in a way that approach you laid out might help us?

>> Yeah, that's a great question, and I appreciate the questions from Kay and Nicol about this.

There's something very interesting in the way we talk about trust with AI tools.

First of all, we talk about AI as a kind of a thing we have to build a trusting relationship with.

And I think I have come to reject that idea entirely. That is not the especially point, that is not what we should be asking for or demanding.

Rather, we should be thinking about who is building these tools, what they are offering us. I don't have a trust relationship with a hammer. Or I hope I don't. I don't have a trust relationship with my microwave. I don't have a trust relationship with my Swiss Army knife. I use them. I use them for specific things. I understand where I can use them. And I don't need to use them for all kinds of stuff. The problem with AI is we have been kind of de-personalizing the active role that companies are playing in constructing a vision of AI, the sledgehammer vision, and presenting that vision to us as the only thing we have available to us.

So it's a tool, it's a thing, we either have to trust it, if we don't trust it we are luddites, if we do trust it we are falling into the hype and that discussion doesn't go anywhere. The reason I think community-driven evaluation is important, I have come to believe this more and more, is that it allows us to figure out, on our own terms, what kind of tools we want to solve what kind of questions and how they can help us.

If you ask healthcare professionals and people working in public health and those going into community to talk to people about vaccinations about what they need, I doubt they would say we want AI to help us. They want help and things to help them in their work. If we try to build things for them, the question of trust would not come up. They would just be using tools they actually wanted to use. That's why community-driven ways of thinking about it is so important. But as a computer scientist who thinks about how to build these technologies and rebuild and reimagine these technologies and is frightened by the hype of companies who want to present one tool to solve our problems, it makes me restless. I'm like that's not something I wanted, make me something different I actually care about. Or as a scientist I can help you build this tool what you want, in a way that is on your terms. That's why I think community-driven evals are so important because it opens that discussion we have been unable, because we are stuck with the binary either you embrace everything the tech companies offer you or you are left with

nothing.

>> KRISHNA UDAYAKUMAR: Thanks very much. Nicol, you want to jump in?

>> NICOL TURNER LEE: I want to respond to that. Suresh and I know each other, I want to push back on something you said about trust, if you don't mind.

>> I don't mind.

>> NICOL TURNER LEE: I trust in my microwave, like you said, or dishwasher. What we trust is it will work like a microwave and a dishwasher, right? And we have a history, a regulatory approaches where we have the Energy Star rating for example, which I write a lot about. Where you go into a big box store you see the big yellow sticker, it says this has been tested by consumers and industry it will give this much water and this much electricity. I think to your point, people want to know there's some level of trust in design. There's some trust like you said in terms of who making it. But more importantly, there is trust it will do what it says it does. I spent a lot of time with facial recognition technology for ex E., like my colleagues and you both. I want to trust when I use it I won't be misidentified as a criminal simply because the technology itself is not optimized to see darker skin hues or mismatch the face and detection part of it. So as a result I now have user technology that has taken away my trust and I can actually put this in my house and look at surveillance of people who might be robbing my house, you know what I mean?

I get what you are saying, I think it's the same in healthcare. We trust that our doctor is going to do the right thing. But we also trust the doctor will use the right scalpel and telescope and other tools that will help us in our recovery. Trust is a huge factor and when you throw in AI and tell people I will take notes on your condition, and they don't know what the notes are. Somebody else gets the notes and they forget the AI transcription was wrong.

When we have those fractures outside we have to demand from the technology, that it's not trust, but trustworthy. That's what we struggle with, how do you make this a trustworthy contextualized technology that people said it did what it was supposed to do and I can actually vouch for that technology.

>> SURESH VENKATASUBRAMANIAN: I mean, you make very good points. I remember being at an event where we were talking about AI regulation and AI governance. The speakers who looked up at the stage and said we have these light bulbs on the stage, we have this building and construction. We trust that the building won't come down on our heads that the lights will soon work. And we trust not because we trust in the electricity or believe in physics but we know there's a regulatory apparatus, there are

standards.

>> NICOL TURNER LEE: That's right.

>> SURESH VENKATASUBRAMANIAN: There's a debate about AI regulation and what we should do and a lot of people especially in the tech industry that we need to stop regulation, not do regulation of AI so we can innovate. I think we forget, or they want us to forget that for every single piece of tech we use in our house, our medicines our drugs, our cars, our consumer appliances they have been rigorously tested over and over again. That's why as Nicol says, we can trust them. Not because we trust in physics or science but we trust in the process, right?

So we need that for AI as well. I have to say right now there's a lot of push to stop doing that for reasons that just totally make me aghast and scared but not based on thinking about the science.

>> KAY FIRTH-BUTTERFIELD: Absolutely. And just to follow-up.

>> KRISHNA UDAYAKUMAR: Kay, I want to come to you. We often see tension on one side we need regulation to protect public safety and ensure trustworthiness. And on the other, the room to innovate. We have seen with AI different approaches with the U.S. than from Europe. What you are seeing is massively more investment into AI and tech development in the U.S., perhaps partially driven by the difference in the regulatory regime that is opening more innovation and perhaps less trustworthiness in the U.S.

Kay, what is the role of government here? If it's not government what are the other ways that this type of trustworthiness can be built? Is it self-regulation? Which hasn't traditionally worked well in other industries. Is it something else?

>> KAY FIRTH-BUTTERFIELD: Thank you for that question. I was going to come in and say actually some of the places we do see where AI is deployed better in the companies that I work with, is in companies that already understand regulation. Say financial services for example. And healthcare.

The willingness to take chances in AI is not as high where somebody understands regulation. And so there is. There is some self-regulation. There's more self-regulation now because most companies, not the AI companies but the companies who use AI are all worried about how they deployed this thing. But now has showing that there are problems with it.

Without government to help them set up frameworks they have to set up their own frameworks which is why I'm not out of a job.

Some of them just use the EU AI Act as a baseline for those trading with Europe and have got to anyway.

The Colorado AI Act actually looks almost identical to the EU AI Act. It's been passed but on hold for a moment once they review it. Some companies are using that. And there are, or there we think there's no legislation in the U.S., actually there are large numbers of states which have legislation and regulation in place. So it's actually for companies it's a bit of a mine field and I trade with this state and that state so how am I going to do this? Which is how they come to saying we are just use the AI Act because it will cover everything and not just human resources or whatever.

I actually don't think it's about whether there's regulation or whether there's no regulation. if there was as much money to invest in Europe. The huge sums of money we are seeing are located amongst businesses based largely on the West Coast.

One thing I will say is that with a lot of companies, and I think particularly in healthcare, if you do one deployment well, it gives you the courage and the knowledge to go on and do another deployment well.

>> SURESH VENKATASUBRAMANIAN: Krishna, can I add a point. Every time I hear the word self-regulation, AI gets its wings and starts fluttering. When I was in the Biden Administration, doing AI policy work and developing the blueprint for AI Bill of Rights we got a lot of questions from companies telling us how they did self-regulation and we didn't need to put out guidelines because they were doing it themselves. We listened to what they were saying and tried to listen where they were coming from. I had many colleagues who worked in these companies. Okay, we should work together and let's see what you can do and how you can regulate. The minute the administration changed, most of these companies completely abandoned their efforts to do any regulation of AI and it was full innovation full speed ahead. The Colorado AI Act Kay mentioned she mentioned it has been on hold for two years. Let me tell you why. Others have been lobbying for two years to kill it. They don't want regulation at all.

And the legislation across the states, there are about 1,000 bills that are in process at different states for the last two years. We have been studying some of this, only Colorado had passed, of all of them. And that's on hold right now. There was a comprehensive bill in California that's come up for review three times and all three times been rejected by the governor or in the legislature. There are smaller bills on AI companions and therapists. Companies do not want AI regulation, even if they say they do. That's why it needs to come from the outside. They are welcome to put in their own standards and do their own best practices. That's great and they should do it. But let's not

pretend we can expect companies to do this on their own. We need something from the outside or they won't comply. I used to not feel this way but I have seen the change quickly happen and I have changed my position on this.

>> KAY FIRTH-BUTTERFIELD: I want to come back on that. I agree entirely, that's why I made the differentiation between the AI companies, producers of AI and the users of AI.

And I agree that everybody should be regulated. But particularly, there is a difference between attitude of those people who are producing AI and those people who are downstream using it.

>> NICOL TURNER LEE: I was going to add on this conversation, which I think is fascinating, particularly as we present it as part of the Bicknell Lecture. I think there's counter positive between innovation and regulation. Some seem to think you can't do both. As Kay said in the European Union, and China, as a matter of fact also has AI regulation, there's this idea that you will slow down and walk cautiously and stop the next big idea.

But as I do in my work, I think the part of regulation that sort of complements innovation is you are controlling the trade-off. Regulators are in the business of protecting the public interest. And so what that means is, they are protecting the public interest around, is this particular AI tool going to have any type of recourse or consequence for people who rely on it in healthcare? We already know in healthcare we have existing statutory requirements and laws. Those don't go away but it makes it clear in the AI space in the use of proxies because you aren't using race or location. You might be using the amount somebody pays for hospitalization, which didn't work well for one company when they used that because they were disproportionately impacting other folks.

The question is how do you get regulators to understand they don't have to know the ingredients of the black box. What they have to know is what is the consequence of the AI on the everyday person? So I think there has to be a conversation shift as well, so you begin to tell people they can actually co-exist.

I also want to say since we are talking about policy, particularly for those on the webinar as well and going back Krishna to this conversation on regulation. I believe self-regulatory frameworks they are more window dressing than substantive. But I think there's something else about being cautious about the government when it comes to AI in healthcare as well.

There's a recent announcement in this administration of the move towards, for example wearables and digital health tools that people can track their health cadence. We have to be

careful because the biggest legislation we lack is privacy legislation. It's not what if I have that wearable and all the sudden I'm not getting my steps in. It's not so much my doctor will be mad at me but so will my health insurance company, and my underwriter and so will these other people. Government medical surveillance is also enacted when we place these tools in the hands of government who hasn't really decided on the regulation that we need to do to protect consumers. Again government plays well in protecting the public interest and that's an area they continue to not go towards. This is the focus of my book, none of this could happen with their patients unless they are connected to the internet. I continuously tell people I'm so excited about AI and possibility of compute power and what we can do, I tried to be Debbie downer for many years. Sometimes I have to come to the other side once in a while because I have seen interesting AI tools who are impacted by health disparities. But the same token it's important to recognize much of what we are asking every day patients to do mean they need to live in areas with high broadband access. Rural communities may not be able to benefit from AI in the ways we think they can. Some urban communities with lack of competition potentially not. Families already distressed in paying services are not near any type of fiber connection, another problem.

So we have to be cautious, I come back to the word ecosystem that we are placing AI in healthcare in a very livable sustainable ecosystem, so we don't place expectations on people they cannot fulfill for us.

>> KRISHNA UDAYAKUMAR: Thank you. And just to keep going down this road a little bit, just a note if there could be a strong role for government and regulation as we have pointed out. But trust in government is also quite low. So thinking about the role of government and if there's not a trustworthy government, which we are seeing all around the world, right? Or governments that are not mature enough in their policy framework, especially for a lot of us that work in lower income countries where the framework doesn't exist, much lessen forced over time. Thinking about other actors and how this could happen in a more collaborative environment, I think, will be really important for how this field moves forward.

But Nicol you mentioned this idea of privacy and you have talked about data before. Let's take a bit of a dive into this side of things. All of this is built on data. So what needs to happen to create a more responsible approach to the data environment and the data infrastructure that then could feed into an AI economy appropriately?

>> NICOL TURNER LEE: You know, I will kind of mention it

just so people know. The AI Equity Lab started workshopping ten different things and my particular interest is in education and healthcare so I spend a lot of time with healthcare people. I find this area to be one of the most consequential inputs we will have because it affects a lot of people. Data is the currency for AI. I know Suresh will hate that I say this word, it's not good data, it's all data companies can get their hands on to train these models. For those less familiar with the artificial intelligence space, it's not just data that allows for predictions, it's also your voice, your face, your text, it's your language that essentially happens. What happens with these models we have been talking about, the 5-6 companies controlling the world right now when it comes to artificial intelligence, they are building large language models which means they have to feed the beast. They have to give more and more data to that. I had the opportunity to hear Karen Howe in Amsterdam last week, she brought up a point I hear, I'm glad you brought up, Krishna, the international context, I hear from the global south, and India and Latin America, small curated data particularly in healthcare where we can account for some of the predispositions to disease that happen close to the equator, versus those that might happen in a different climate when it comes to people or potential patients or people who might be subjected to certain disease.

I think that's an interesting way to look at it and debunks the model that all data has to be big, massive forms of collection, et cetera. Another thing happening that is important for this community to know as well. We are finding and I like how Kay is posing this with the business community, we are finding small startups and innovators that are beginning to look at ways in which you direct data that makes sense for different sectors.

So I literally just sat on a panel with a woman who worked at the Gates Foundation, she is a Fellow, she is working with a company working with healthcare data ensuring its inclusive and includes different types of predispositions, it's something they can authenticate and validate and audit.

These are small companies that are essentially feeding larger companies or universities with regards to that. I think, Krishna, to your point, there are lots of challenges in collecting data, based on who is collecting it. Lots of challenges having it representative.

I think we have to be careful in the health space to recreate the same types of models with clinical trials where we compensate people for their data. I think that's one of those areas I sort of cringe at. In our rush to get data we also think maybe we should over sample and include more people and give

them \$50, with some of these small data companies do. But when you start to get into health, hospitals, large medical institutions were to do that, there are lots of liabilities and risks that come with that. In addition to the fact that you are now incentivizing people to give up something that is personal to them without the appropriate guardrails.

So this kind of goes back to the regulatory question. There's regulation at government, but there should be regulation at your institutions.

There has to be as Kay said, soft and hard guidance. Red and green lines which allows you to work externally by asking the right questions. Someone in the chat said what questions do we need to ask? It allows you as a hospital if you want to collect consumer data to build an AI model or system you also have to do the research and grunt work to do continuous oversight to make sure people will not be harmed.

So at the end of the day, for me, it's a question that we are not going to be able to solve. I have been trying to figure this out with a lot of people. How do you solve the fact there's no more data, I wake up from a hot sweat thinking about Henrietta lax and I say it starts with, again, Suresh said participatory engagement so you actually get the right people at the table to come up with the care, knowledge and human oversight. So we all got a lot of work to do and all of you on this call as well. It's worth pointing out, I have been noticing in this data conversation, it's beginning to look more smaller than larger when it comes to the implementation of models.

>> KRISHNA UDAYAKUMAR: Thank you. Kay, let me bring you into this conversation about data and Suresh as well. Maybe add to this context that the good actor, so to speak are the ones being left behind, because there's a lack of strong rules and regulations. If you want to do something, you can.

We are seeing misinformation spread wildly while those of us trying to curate good information around health and public health are doing it in very small ways.

How do we account for strengthening the data environment in a place where there are not guardrails for all actors who want to move as quickly as they would otherwise?

Kay, your thoughts on that?

>> KAY FIRTH-BUTTERFIELD: It's really hard. And I would say almost impossible. I see my fellow panelists agree. I just want to, I think to add onto Nicol's point about access to the internet. When we think about serving people in the global south, 2.6 billion people are still not connected to the internet.

So you know, I hear people wildly saying, oh we will be able to use AI to stop women dying in child birth. No you won't.

Because the majority of women dying in child birth, there is no connectivity.

So one of the big worries for me, as we think about public health is that split. We are seeing the split within the United States and other countries between those who have and good access to the internet. And Nicol talked about that earlier and also internationally as well.

I think that we absolutely have to have good data protection rules and we don't and until we do we won't cure the misinformation and we can see the results with the measles epidemic.

>> KRISHNA UDAYAKUMAR: Suresh?

>> SURESH VENKATASUBRAMANIAN: Yeah, it's tricky because I think as we are seeing this conversation go from AI to data to the global landscape they start to feel increasingly interactable. I think it could be frustrating to think how all these things are connected in one gigantic hair ball that we can't disentangle. I would like to offer maybe a hopeful perspective. Let's go back to Nicol's point, small data, excellent point about the value of small data.

There was a post yesterday, ten years ago in internet land but yesterday on a new tool for building a small language model. Which is like 8,000 lines of python code, maybe training for about \$100 of compute time on a node can rent from AWS. You get something that is reasonable. This isn't something we expected to do two or three years ago. The fact we can do that now, it's possible to do that now means there is innovation. You will use the word innovation in building smaller models that use less resources or more target resources. Another point. I'm hearing more and more from people the following, yeah, I use ChatGPT, it didn't really give me the answer I wanted. But this company build this specific customized model for my use case and that seems to work very well for me. Even though all these chat bots don't work very well. I find that encouraging people are thinking to build specific boutique, artisanal models, if you will, you could actually see benefits without having to amass everybody's data and do all kinds of surveillance.

I say that because A, it's happening and it's an innovation of a real kind and something we can ask for and demand. We don't have to take for granted the only tools available are the ones that require web scale, data scale that have to be western oriented or English because that's all the data we have. We can demand more and ask more and this will be companies and researcher that's will give it to us. That's where I feel hopeful, by people saying no, I want better, I want a better model, I want better tools that do what I want, somebody will be out there to give it to you. Maybe a 15-year-old looking for a

new project who will build that for you. But they can do it. That's where the real innovation and new ideas will come from. That gets us out of the trap we seem to be in with data and AI and so on.

>> KRISHNA UDAYAKUMAR: Beautiful. Anybody want to add onto that?

>> NICOL TURNER LEE: So you just had me think of this conversation I had recently about criminal justice data. You can probably take healthcare and put in criminal justice, financial services, we were talking about small language models. Wow this person did 20 years of work with the justice impacted community. Do you know how much great information I have from when people are incarcerated to when they leave. She said this word, she said I have values-driven data.

Meaning, when I think about this data, I give it this care because I am close to this community in so many ways, personally, professionally. And my data reflects questions that would never be asked for someone who doesn't know that.

It goes back to what you said, this participatory model, I did the Equity Lab, we are kind of the OG's who have been watching this thing for quite some time. This question of values is one in which has not fit, squarely in the responsible trustworthy, ethical naming and narrative that we put out about AI. And I think that's where, again, people are becoming much more reasonable to question this technology and to say, do I really need this? And I think Kay's point, if we can go back to the Character AI, we are now starting to see groups resist the fact that we have to live in this world, where we cannot be luddites. I think that is a question with any type of technology, this is not new. We have seen this over the course of generations, when a new technological revolution happens in our country. But it also suggests when we find ourselves stuck, do we not only know we have agency to say no, but do we also know how to get out of it. That's the only thing I would say to my colleagues. I don't know if they want to take this on. Character AI exists because we live in a lonely society.

It doesn't exist because people found more empathy and partnership. It's because we live in a society where people feel disconnected, unloved and alone. And so, technology in many respects is building off our fragility. And that's a question for those of us in the health space, have to continue to think about. Is this a tool that is going to help me be a better physician, or a tool that is filling a blank of something that I need to be doing in my role, you know, to enhance the human experience? And that's something I think we need to teach medical students, early on, who are enticed by going into AI and using it in these spaces that require the human experience.

>> KAY FIRTH-BUTTERFIELD: Yeah, I will just add onto that, I mentioned my breast cancer experience. When my oncologist found out what I did, she said, oh wouldn't it be amazing if we could have AI help you with your journey with breast cancer?

And I looked at her and I said I think it would be better if we could get AI to help you with some of the tasks you have in your back office so you, the human being, can walk me through my relationship with cancer.

And my cancer journey.

So, for me, it all comes back to that first light of mine, humanity.

We are told we are in the age of AI. I would like us to think we are in the age of humans and we have this tool called AI that may be able to help us, if we use it wisely.

>> KRISHNA UDAYAKUMAR: That's wonderful. We are coming up to the end of our time. It's been such an amazing conversation. We have gotten so many great conversations and comments and we incorporated as many as we could into this discussion as well. Let me ask you for a final thought you would leave our audience with and maybe put some focus on it. We are here with the School of Public Health. Many you work with students and young people all the time. As they are thinking of their careers in an AI economy, their role as students.

What's a message you would send to them about what's most important to hold dear, what's most important in terms of the opportunities as they think about what the future actually holds here? Suresh, I'm going to start with you.

>> SURESH VENKATASUBRAMANIAN: The thing I always tell people which comes as a surprise, these are fun tools to play with. You should just embrace them, just play with them. Not because you are either giving into the hype or sucker for the technology or whatever. Just because they are fun to play with.

In doing so, you will learn, to have your own relationship, if you wish. Your own way of interacting with different kinds of tools, seeing where their warts are, seeing where they are helpful. There are times, I use these tools a lot for example, writing small snippets of python code for graphs to settle bets with my friends about sports statistics. It works well for me and I win the argument. But anyway, I have learned over time how I can get them to work and where they don't work as well.

I have built that understanding. We may build understanding with many technology with search engines or maps or whatever. I think people shouldn't feel at all hesitant to do that. That is separate, that will give you a much more informed perspective, in your own domain whether they make sense to be used in your domain or not. Otherwise we are stuck at still having conversations at a level which is not informed by actual lived

experience by the tools. We could have that and have that more informed conversation. So that's what I would say.

>> KRISHNA UDAYAKUMAR: Thank you so much. Nicol?

>> NICOL TURNER LEE: Yeah, I agree with Suresh. Although I don't play with the tools because honestly I don't know how to use them, so. Both of my kids are in college and graduate school, just me and the dogs and they can't seem to teach me. I have this experiment, you will never touch a chat bot, then I was trying to get a book at Barnes & Noble, and I used it, I use AI when I'm actually looking for something very discrete. I think goes to Suresh and Kay's point, use it the way you see fit. I see people putting in the chat and this came up on my panel last week, we need a frame for how we are teaching medical students, we need one track on the productivity and efficiency and that track, there was an article in The New York Times that said the AI transcription of a medical patient still ended up with him as a fatality because there was a misread of the notes based on the AI transcription. We need to teach these tools exist but we need them to be great interrogators. I'm not a practicing doctor but Ph.D. in sociology, there are sub sectors. So I think gastroenterologists need something different than heart health. They should go to the American Heart Association and say we need guidance. And here is some things we think need to be in this broad stroke guidance on AI's application in our field. You know, when it comes to thinking about the condition of the heart. Thinking about the radiology imaging, thinking about those kinds of things and we need to do the same thing, Krishna, with public health.

I think this needs to be a bottoms-up conversation that disrupts power. If we do that we might get further to have more agency, but when it's all said and done we need to teach medical students if you don't know how to use it and you still want to take that telescope and carry it around and put it on someone's chest do so. We are sending the notice this is the way of the future of medicine. It may be. Just like we didn't know we would have digital health records but we still do it in a thoughtful and caring manner within our sectors themselves.

>> KRISHNA UDAYAKUMAR: Thank you. Kay?

>> KAY FIRTH-BUTTERFIELD: I agree with Suresh. But use it wisely, I would say. Because there are true climate problems here at the moment. As you are using it.

I think that as humans, the one thing we need to do is think critically about these tools, both the way we use them in our homes and in our own lives and in our practices.

And you know, use your own intelligence to think about whether this particular tool is going to make your practice better for you and perhaps more importantly better for your

patient. Is it going to improve patient outcomes? If it's not, then probably don't use it.

>> KRISHNA UDAYAKUMAR: Yeah. Well, fantastic. We could keep this going all afternoon. But unfortunately we are out of time.

What an amazing tour de force we have had from the three of you. We really started a conversation around technology and ended the conversation around people and humanity. So often when we see technologies that have transformational potential, there's this tension that we have all talked through of how do we use it for the best possible ways, while we are managing the risks. So how do we think about the interaction between technology and people as the path forward and make sure we are creating guide posts, everything from the right evidence and ways to generate evidence. The ethical and responsible development and use. The issues around bias and transparency. Certainly around sustainability of an AI economy going forward and what does that mean for climate. And also jobs in the future.

And of course equity and fairness and access to these tools as they develop over time. We heard about focus on use cases, making sure we are developing better tools for important cases as ways to improve health, healthcare and public health going forward. And how this could be grounded in the community and values that should really be driving the path forward in the field.

So let me first thank Nicol, Suresh, Kay for your expertise, for your time. Such fantastic perspectives to make us all a little smarter through this. A huge thank you to the BU School of Public Health, data science, hopefully we honored the legacy of William Bicknell with this conversation. I will turn it back to Dr. Debbie Cheng to close us out.

>> DEBBIE CHENG: Thank you to our moderator, Dr. Udayakumar, to our speakers and everyone for joining us today. We hope Dr. Bicknell's vision continues to guide and resonate as we explore bold ideas in public health.

Please be sure to register and join our next and final online conversation of the semester, public health and new media. Modes of persuasion. Co-hosted by the BU College of Education and Public Health Post. November 12th at 1:00 p.m. eastern.

Thank you.