

**ARTIFICIAL INTELLIGENCE AND HEALTH
CARE: PROMISE AND PITFALLS**

HEARING
BEFORE THE
COMMITTEE ON FINANCE
UNITED STATES SENATE
ONE HUNDRED EIGHTEENTH CONGRESS
SECOND SESSION

FEBRUARY 8, 2024



Printed for the use of the Committee on Finance

U.S. GOVERNMENT PUBLISHING OFFICE

62-427—PDF

WASHINGTON : 2026

COMMITTEE ON FINANCE

RON WYDEN, Oregon, *Chairman*

DEBBIE STABENOW, Michigan	MIKE CRAPO, Idaho
MARIA CANTWELL, Washington	CHUCK GRASSLEY, Iowa
ROBERT MENENDEZ, New Jersey	JOHN CORNYN, Texas
THOMAS R. CARPER, Delaware	JOHN THUNE, South Dakota
BENJAMIN L. CARDIN, Maryland	TIM SCOTT, South Carolina
SHERROD BROWN, Ohio	BILL CASSIDY, Louisiana
MICHAEL F. BENNET, Colorado	JAMES LANKFORD, Oklahoma
ROBERT P. CASEY, JR., Pennsylvania	STEVE DAINES, Montana
MARK R. WARNER, Virginia	TODD YOUNG, Indiana
SHELDON WHITEHOUSE, Rhode Island	JOHN BARRASSO, Wyoming
MAGGIE HASSAN, New Hampshire	RON JOHNSON, Wisconsin
CATHERINE CORTEZ MASTO, Nevada	THOM TILLIS, North Carolina
ELIZABETH WARREN, Massachusetts	MARSHA BLACKBURN, Tennessee

JOSHUA SHEINKMAN, *Staff Director*
GREGG RICHARD, *Republican Staff Director*

CONTENTS

OPENING STATEMENTS

	Page
Wyden, Hon. Ron, a U.S. Senator from Oregon, chairman, Committee on Finance	1

WITNESSES

Shen, Peter, head of digital and automation for North America, Siemens Healthineers, Washington, DC	5
Sendak, Mark, M.D., MPP, co-lead, Health AI Partnership, Durham, NC	7
Mello, Michelle M., JD, Ph.D., professor of health policy and law, Stanford University, Stanford, CA	9
Obermeyer, Ziad, M.D., associate professor and Blue Cross of California distinguished professor, University of California, Berkeley, Berkeley, CA	11
Baicker, Katherine, Ph.D., provost, University of Chicago, Chicago, IL	14

ALPHABETICAL LISTING AND APPENDIX MATERIAL

Baicker, Katherine, Ph.D.:	
Testimony	14
Prepared statement	41
Responses to questions from committee members	43
Crapo, Hon. Mike:	
Prepared statement	45
Mello, Michelle M., JD, Ph.D.:	
Testimony	9
Prepared statement	46
Responses to questions from committee members	48
Obermeyer, Ziad, M.D.:	
Testimony	11
Prepared statement	55
Responses to questions from committee members	58
Sendak, Mark, M.D., MPP:	
Testimony	7
Prepared statement	60
Responses to questions from committee members	62
Shen, Peter:	
Testimony	5
Prepared statement	69
Responses to questions from committee members	76
Wyden, Hon. Ron:	
Opening statement	1
Prepared statement	80

COMMUNICATIONS

AARP	83
Advanced Medical Technology Association (AdvaMed) Imaging	85
AHIP	86
American Federation of Teachers	90
American Institute for Medical and Biological Engineering	92
American Medical Association	93
Asher Informatics PBC	101
Center for AI and Digital Policy	103
Connected Health Initiative	107

IV

	Page
Federation of American Hospitals	118
Healthcare Confidentiality Coalition	119
Healthcare Leadership Council	120
Medical Group Management Association	123
National Health Council	125
National Health Law Program	127
National Nurses United	130

ARTIFICIAL INTELLIGENCE AND HEALTH CARE: PROMISE AND PITFALLS

THURSDAY, FEBRUARY 8, 2024

U.S. SENATE,
COMMITTEE ON FINANCE,
Washington, DC.

The hearing was convened, pursuant to notice, at 10:08 a.m., in Room SD-215, Dirksen Senate Office Building, Hon. Ron Wyden (chairman of the committee) presiding.

Present: Senators Menendez, Carper, Cardin, Bennet, Warner, Whitehouse, Cortez Masto, Warren, Thune, Cassidy, Young, Johnson, and Blackburn.

Also present: Democratic staff: Melissa Dickerson, Senior Investigator; Eva DuGoff, Senior Health Advisor; Marielle Kress, Senior Health Advisor; Marisa Salemme, Senior Health Advisor; and Joshua Sheinkman, Staff Director. Republican staff: Gable Brady, Senior Health Policy Advisor; Kellie McConnell, Health Policy Director; Gregg Richard, Staff Director; and Conor Sheehey, Senior Health Policy Advisor.

OPENING STATEMENT OF HON. RON WYDEN, A U.S. SENATOR FROM OREGON, CHAIRMAN, COMMITTEE ON FINANCE

The CHAIRMAN. The Finance Committee will come to order. The first thing I want to say to our guests is, obviously this is a very hectic day in the U.S. Senate—something of an understatement—and I want our witnesses to know that our colleagues are all trying to juggle responsibilities.

So I do not want our witnesses to feel in any way that the fact that Senators will be coming and going minimizes the importance of this hearing. And my friend and partner Senator Crapo is an example of trying to be several places at once. And so, as we begin, I want to ask unanimous consent that Senator Crapo's prepared statement be entered into the record after my opening statement.

[The prepared statement of Senator Crapo appears in the appendix.]

The CHAIRMAN. This morning, the Finance Committee meets to discuss the use of artificial intelligence in health care. The focus is going to be on the technology that's being used in Federal health programs such as Medicare and Medicaid. There is no doubt that some of this technology is already making our health-care system more efficient.

But some of these big data systems are riddled with biases that discriminate against patients based on race, gender, sexual orientation, and disability. It is very clear in my judgment—and technol-

ogy is an area I have tried to specialize in since my arrival in the U.S. Senate, when only Senator Pat Leahy knew how to use a computer—it is very clear that not enough is being done to protect patients from bias in AI.

We work to ensure innovation. For example, in the 1990s we improved patient care, and we empowered telemedicine, digital signatures, and other efforts. Congress now has an obligation to ensure the good outcomes from AI set the rules of the road for new innovations in American health care.

Today we are going to discuss the role Congress and the committee must play in helping strike a balance between protecting innovation and protecting patients and their privacy with legislative proposals like the Algorithm Accountability Act, which I have introduced with my colleague and friend Senator Booker, and Congresswoman Yvette Clarke, a very knowledgeable member of Congress on technology issues. Our legislation would tackle these concerns head-on.

There are a lot of reasons to be optimistic about the potential of AI to improve health care. Today, the industry faces a host of challenges, all made worse by the strain of the COVID pandemic on our health system. There is a workforce shortage; existing providers are facing high rates of burnout; health-care costs are rising faster than wages; and there is an ever-growing gap between the care that is needed and the care that is being delivered.

Already, AI tools are being deployed to reduce some of these pressures and ease the strain on the industry and providers. Some doctors use the technology to prepopulate, for example, clinical notes and their emails to reduce workloads, submit bills to insurers, and help to reduce administrative waste and even help with diagnostics.

Primary care providers can use these tools to screen for certain diseases and connect patients with specialists for treatment that saves patients time and money and leads to better, more timely care. So, that brings us to the area that this committee has had a special interest in, and that's Medicare and Medicaid.

Here is an opportunity to improve workload for providers and help patients, all of whom are trying to make sense out of this new AI reality. And addressing these challenges with new technology has to mean better patient outcomes, while at the same time protecting privacy.

And that goes right to the heart of my philosophy, for now several decades, with respect to technology. Technology gives us a chance to innovate, and that innovation is not mutually exclusive when it comes to privacy. Smart policies give you both. They give you innovation and privacy. Not-so-smart policies give you less of both. So, as we begin this effort with respect to AI and these developments, let us keep that in mind.

So, there are clear, glaring examples of AI tools being developed with data that perpetuate racial biases, and I have been pleased to be working on this with Senator Booker and Congresswoman Clarke, who have really zeroed in for all of us in both the Senate and the House on some of these issues, because these biases have been deployed in ways that bypass important doctor expertise, and that leads to inadequate care for patients.

So the committee is very lucky today to have Dr. Ziad Obermeyer, who in 2019 discovered racial bias in an AI tool developed by the health-care company Optum, a subsidiary of the United-Health Group, that was used by providers across the country to offer care management services.

Dr. Obermeyer found that the tool on average required Black patients to present with worse symptoms than White patients in order to qualify for the same level of care. Folks, that is not a close call. It is just not! And we have seen it in so many other areas, just in the last few days still trying to sort through the concussion settlements that have been discussed with respect to NFL players. So I am very pleased that Dr. Obermeyer is here, and we appreciate his expertise.

That algorithm was available to thousands of doctors across the country, potentially impacting millions of patients. How does such a flawed system make its way into general use? Well, it's not very hard to figure that out. Nobody is home. Nobody is watching. No guard rails. No guard rails to protect the patients from flawed algorithms in AI systems.

To make matters worse, the technology the insurance companies or health systems use can play a role in what care patients receive—and what services are approved or denied. The Department of Health and Human Services does not, as of today, really oversee the use of these systems. Big problem, folks.

Most of us here would agree that there are many ways this technology can be used to improve health care and patient outcomes. As long as we increasingly rely on technology like AI to make decisions in every part of our day-to-day lives, the Finance Committee—we are going to work on a bipartisan basis.

Senator Crapo and I have talked about this a number of times. We are going to work in a bipartisan way to deal with these crucial issues. I happen to believe that one of those keys is to have guard rails in place to protect patients, particularly in Medicare and Medicaid, and I do not believe the current laws go far enough to deal with that.

That is why we came forward with the Algorithm Accountability Act. It does not answer all the questions, but it is all about common sense. We talked to technologists, we talked to authorities, particularly about some of the first steps, and that is what we did with the Algorithm Accountability Act to lay the groundwork to root out the algorithmic biases from these systems.

As applied to health care, our legislation would require health-care systems to regularly assess whether the AI tools they develop or select are being used as intended and are not in effect generating more, and what amounts to perpetual, harmful bias.

I will close with this. The same protections in my Algorithm Accountability Act have to be in Medicare and Medicaid. So what we need, if I could sum it up, is transparency in how the tools are developed and used to foster trust and accountability for how they are used in health care, making sure we preserve the privacy of patients, and letting us use this as an opportunity to give everybody in America the chance to get ahead.

I mean, it is really about equity. That is what I want to have our committee work towards on a bipartisan basis. When you look at

what was done in this committee room with the historic tax reform bill—it was in 1986 before a lot of our audience was born—it was all about giving Americans, everybody, the opportunity to get ahead.

And that is our country at its best. That is what we are all about, and we have to make sure these tools further equity in health care and do not perpetuate harmful bias or disadvantage hospitals and providers who service low-income patients or communities of color.

The Food and Drug Administration and the Office of the National Coordinator for Health IT have proposed some new rules. I think they are a step forward. My own take, speaking for myself—my colleagues, as you know, are juggling a lot today—I do not think these rules go far enough.

I believe more is needed to protect patients from flawed systems that can and will directly affect the health care they receive. I look forward to working with my colleagues on the committee to identify ways we can protect patients and improve their care going forward.

I say to our guests, that is what this committee is all about: working in a bipartisan way. We did it this Congress with PBMs, for example, these middlemen. We have done it with respect to mental health.

We have a lot of issues on our plate, but we try to tackle them in a bipartisan way. And I want to thank our witnesses for testifying at today's hearing, and I look forward to hearing them.

[The prepared statement of Chairman Wyden appears in the appendix.]

THE CHAIRMAN. Now, let's see. We are going to have to introduce these wonderful people. Peter Shen is here, director of digital and automation, North America at Siemens. He focuses on introducing new and emerging technologies in the health-care field. He is an academic star in biomedical engineering and mathematical sciences from Johns Hopkins.

Mark Sendak is co-lead at the Health AI Partnership. That is a collaboration between academic systems and businesses and Federal entities. He is the population, health, and science lead at the Duke Institute for Health Innovation. If I go on and on about Dr. Sendak, you will be here till breakfast. But we are glad he is here.

Michelle Mello is here, and Michelle comes from my alma mater. You know, I wanted to play in the NBA. I got a scholarship to Cal at Santa Barbara, and I did not get to Stanford as an undergraduate until it was clear I was not going to make it. So glad you are here, Dr. Mello. She leads empirical health. She's the empirical health law scholar, and her research focuses on understanding the effects of law and regulation. As I say, she is a professor now both at Stanford Law School and Stanford University School of Medicine. Stanford Law School is no longer right across from ugly anymore. It has changed. Glad you are here.

Dr. Obermeyer—I threw some bouquets out already. We are just really pleased he is here. And he is an associate professor and the Blue Cross distinguished professor at UC Berkeley. He was named one of the 100 most influential people in AI by *Time* magazine for his work that I discussed already.

And, Doctor, am I pronouncing this right: Biker?

Dr. BAICKER. Baker spelled funny.

The CHAIRMAN. Baicker. Yes, that is right; an Oregonian, great. Leading scholar on the economic analysis of health-care policy. She served as a Senate-confirmed member of the President's Council of Economic Advisors. She received her Ph.D. in economics from Harvard.

We welcome all of you, and we will include all your prepared remarks as a part of the record. Dr. Baicker, where in Oregon are you from?

Dr. BAICKER. I am unfortunately not from Oregon.

The CHAIRMAN. Oh; I thought I heard you say you were from Oregon.

Dr. BAICKER. No. I did a research project based in Oregon and had an opportunity to spend some time there, and had a wonderful introduction to——

The CHAIRMAN. Well, come back. I am glad we have sorted out your connection.

Dr. BAICKER. Well, I now consider myself invited to return.

The CHAIRMAN. You do. All right. We have sorted that out.

Let's go. Mr. Shen?

STATEMENT OF PETER SHEN, HEAD OF DIGITAL AND AUTOMATION FOR NORTH AMERICA, SIEMENS HEALTHINEERS, WASHINGTON, DC

Mr. SHEN. Chairman Wyden, on behalf of Siemens Healthineers and our nearly 17,000 employees in the U.S. and approximately 71,000 employees globally, thank you for the opportunity today to testify on the topic of artificial intelligence in health care.

My name is Peter Shen, and I am the North America head of digital and automation for Siemens Healthineers. My career focus is on the introduction of new and emerging technologies in the health-care market, including artificial intelligence.

Siemens Healthineers is a leading medical technology company with more than 120 years of history and experience bringing breakthrough innovations to the market that enable health-care professionals to deliver the best care for patients. Our core portfolio includes imaging, diagnostics, and therapies augmented by digital technologies and AI. We partner with more than 90 percent of leading health-care providers to address population growth and chronic disease prevalence, health-care workforce shortages, and the lack of access to care to underserved areas.

We have the distinction of being the only medical technology company capable of end-to-end cancer care, from diagnosis and screening to treatment and survivorship. This is a responsibility we take very seriously, as we keep patients at the center of everything we do.

Siemens Healthineers has been working on applying artificial intelligence in medical technology for more than 20 years. At our AI Office of Big Data here in the U.S., we have created and maintained one of the most powerful supercomputing infrastructures dedicated to developing AI. This allows our research scientists to collect, prepare, organize, and secure the identified data needed to train and deliver accurate AI algorithms. From its inception, we created and maintained a transparent quality assurance process,

which involves clinical validation to guarantee the data being used to train the AI algorithms is accurate for diagnosis and treating disease, and is based on a balanced cohort of people of different ages, genders, and ethnicities, thus ensuring that we develop reliable, accurate, and unbiased AI algorithms that protect the patient and are reflective of the patient population that they will be applied toward.

Unlike AI for operational work or workflow improvements that help reduce position burden or improve patient experience, many Siemens Healthineers clinical AI algorithms can be termed as Algorithm-Based Health-care Services. These analytical services delivered by FDA-cleared devices use AI and machine learning to produce quantitative and qualitative clinical outputs for physicians to use in the diagnosis or treatment of disease that were previously not possible to visualize or calculate without the assistance of AI.

The patient journey is at the heart of Siemens Healthineers AI work, and AI is already helping to improve care and outcomes for the patient. Clinical AI or Algorithm-Based Health-care Services can be an important service when used to diagnose neurodegenerative diseases for patients. Changes in brain volume over time can be a powerful predictor of diseases such as Alzheimer's. But neurologists are challenged with needing actionable patient-specific brain volume data to diagnose and treat such patients more accurately.

Our clinical AI algorithm services can automatically segment different structures of the brain on an MRI image, measure their volumes, and compare these volumes to a normal brain database. Brain volume deviations from a norm are highlighted in a comparative report from the neurologist to provide additional, objective, quantitative information that they can use to make a more accurate and informed diagnostic and treatment, resulting in better patient outcomes.

While CMS has recognized the value and the complex nature of Algorithm-Based Health-care Services, the agency's reimbursement decisions have not uniformly and consistently ensured appropriate levels of payment for these services. This inconsistent, unpredictable approach stifles adoption by providers, especially in rural and underserved areas, and therefore restricts patients' access to new and innovative diagnostic tests and treatments.

We support a solution that ensures a predictable and consistent approach by CMS through a temporary and separate payment for 5 years based on manufacturer-supplied cost data. This approach recognizes the cost to develop and integrate AI into the clinical setting and reimburses for the distinct clinical value Algorithm-Based Health-care Services provide.

Guaranteeing a consistent reimbursement process would empower hospitals and providers to invest in AI confidently, ensuring their services are appropriately reimbursed. Without this financial support, these providers will face difficulties in embracing and integrating AI technologies, ultimately potentially denying revolutionary services to patients.

Siemens Healthineers believe AI has the greatest potential to improve access to care, assist physicians in the diagnosis of disease, and enable more personalized treatments for the patient. As a mar-

ket leader in research and training, AI, and medical technologies, we are excited about what the future holds. It is critical that we all work together to ensure we create trust with consumers and build ethical, transparent, and accessible AI in health care to ultimately improve patient outcomes.

Again, thank you for the opportunity to testify before you today, and I look forward to your questions.

[The prepared statement of Mr. Shen appears in the appendix.]

The CHAIRMAN. Thanks for getting us off to a good start.

Dr. Sendak?

**STATEMENT OF MARK SENDAK, M.D., MPP, CO-LEAD,
HEALTH AI PARTNERSHIP, DURHAM, NC**

Dr. SENDAK. Chairman Wyden, Ranking Member Crapo, and members of the committee, my name is Mark Sendak, and I appreciate the opportunity to serve on the panel today. I must note that any views expressed in my testimony are my own and may not reflect those of my employer or the multi-institutional partnership I help lead.

I serve as the population health and data science lead at the Duke Institute for Health Innovation, DIHI for short, and I am the co-lead for Health AI Partnership. I have been developing and implementing AI technologies in clinical care for over a decade.

Since DIHI's founding in 2013, our team has developed and implemented over 20 AI technologies in clinical care. We were the first in the U.S. to implement a deep learning model in routine care. We were the first to implement Model Facts labels for AI tools, and we have incubated four companies to commercialize AI products built at Duke.

Our team has demonstrated the benefits of AI in health care. Duke dramatically improved the quality of sepsis care using our Sepsis Watch system. Duke proactively manages chronic diseases in Medicare patients by using AI to identify patients at risk of complications.

But my comments today will not focus on the amazing work I have been a part of at Duke. Today, I am speaking with you primarily as the co-lead of Health AI Partnership. In 2018, a mentor of mine asked me, "How do we get AI out of the ivory tower?" At that time, my experience with AI at Duke was unimaginable to people outside of a few exceptional islands of excellence, and there was minimal infrastructure being built to get AI outside into low-resource settings.

In 2021, I helped launch Health AI Partnership to advance the safe, effective, and equitable use of AI in all health-care organizations. We exist to get AI out of the ivory tower.

The Senate Finance Committee can take concrete action to advance accountability, equity, privacy, and transparency in the use of AI in health care. The Medicare program ensures high-quality care for beneficiaries through conditions of participation and other mechanisms. There is a unique opportunity for this committee to strengthen Medicare controls on the use of AI, and to facilitate investments in technical assistance, technical infrastructure, and training.

First, we can talk about guard rails. Through Health AI Partnership, we work with 20 organizations across the U.S. to surface and disseminate AI best practices. We interview leaders and run case-based workshops to develop practical resources for health-care leaders asking basic questions. How do I evaluate different externally built AI products? How do I navigate the new FDA clinical decision support guidance? How do I assess the potential future impact of this AI product on health inequities? How do I align organizational processes with the White House blueprint for an AI bill of rights?

Health AI Partnership resources and programs provide guard rails for high-resource organizations that are rapidly accelerating their use of AI. Adoption of these guard rails by hospitals could be required by Medicare program participation, but guard rails only serve the few organizations that are already on the AI adoption highway.

We must also address the more critical need for roads, on ramps, and bridges—the core infrastructure investments needed to ensure that all people in the U.S. benefit from AI in health care. Most health-care organizations in the U.S. need an on ramp to the AI adoption highway. They are struggling with clinician burnout. They face razor-thin or negative margins. They are entirely dependent on external EHR vendors for technology expertise and assistance.

Simply put, they do not have the resources, personnel, or technical infrastructure to embrace guard rails for the AI adoption highway. Core infrastructure investments are needed for technical assistance, technology infrastructure, and training. Fifteen years ago, Congress funded the procurement of EHRs along with 62 regional extension centers to support EHR implementation in low-resource settings.

While EHRs are far from perfect, Federal programs did successfully enable broad adoption of the technology. Health-care organizations urgently need technology infrastructure that is distinct from EHRs that enables the efficient evaluation, clinical integration, and monitoring of AI tools.

Last, training programs are needed to rapidly equip health-care leaders with the foundational knowledge required to locally govern AI. Local AI governance needs to be a core competency for health-care organizations.

Thank you again for the opportunity, and I look forward to answering your questions.

[The prepared statement of Dr. Sendak appears in the appendix.]

The CHAIRMAN. Dr. Sendak—and, Dr. Mello, I am sure you heard your colleague, your seat mate, talk about getting everything out of the ivory tower, and no pressure, but I am sympathetic. I was a guide at Hoover when I was on campus, and I always got lost trying to remember exactly how many books there were, because we were not keeping up.

So, no pressure, but get us out of the ivory tower, please.

**STATEMENT OF MICHELLE M. MELLO, JD, Ph.D., PROFESSOR
OF HEALTH POLICY AND LAW, STANFORD UNIVERSITY,
STANFORD, CA**

Dr. MELLO. Thank you, Mr. Chairman. So you were literally in an ivory tower?

The CHAIRMAN. I was.

Dr. MELLO. Yes. Well, I am so pleased to have the opportunity to speak with you, and I am sorry that I am coming to you today with a voice that is more suitable for jazz singing than testifying.

I am part of a group of ethicists and data scientists and physicians at Stanford that evaluates AI deployments that are proposed for use in Stanford health-care facilities, which care for over a million patients per year. I would like to share with you the three most important things that we have learned in doing this work.

First, while hospitals increasingly recognize the need to vet AI tools before use, most health-care organizations do not have robust review processes yet. They need help, and there are many things that Congress can do to help.

Second, to be effective, governance cannot focus only on the algorithm. It also has to encompass how the algorithm is incorporated into clinical workflows, and by workflows, I mean how physicians and nurses and other staff interact with each other, with the AI tool, with patients, and with other systems.

Conversations about regulating AI mostly focus on the algorithms, but equally important is asking questions about how medical professionals will interact with them. For example, we have looked at the onus on medical professionals to evaluate whether model output is accurate, given the information they have at hand and the time they have available.

Large language models like ChatGPT are used to summarize clinic visits, doctor's notes, and even draft replies to patient's emails. Developers assume the doctors will carefully edit those drafts before they are submitted. But will they?

To address such issues, oversight must go beyond the algorithm to reach how adopting organizations will use it. To take a simple analogy, if we want to prevent motor vehicle accidents, it is not enough to just set design standards for cars. Road safety features, driver's licensing requirements, and rules of the road also help keep people safe.

Third, because the success of AI tools depends on the organization's ability to support them during use, the Federal Government should establish standards for organizational readiness and responsibility to use health-care AI tools, as well as for the tools themselves.

But it would be a mistake to enshrine in legislation detailed standards for such a fast-moving field. We must have the humility to acknowledge we do not know what the right standards will be 2 years from now. Regulation needs to be adaptable, or else it will risk irrelevance—or worse, chilling innovation without producing countervailing benefits.

The wisest course is for the Federal Government to foster a consensus-building process that brings experts together to create standards and processes for evaluating proposed uses of health-

care AI. It can also require that entities regulated by Federal agencies adhere to those same standards and processes.

For example, the Medicare certification process could be used to require that hospitals have a plan for vetting AI tools before deployment and monitoring them afterwards. The initiative currently underway to create a network of AI assurance labs is right on track. These centers can develop consensus-based standards and perform some evaluations of AI tools for organizations that lack the resources to do it themselves. Adequate funding is critical to their success.

Some aspects of AI review have to happen within individual health-care organizations, though. They can best identify problems that arise from integration into workflow. Regulatory requirements can ensure that organizations invest in making that happen, just as the Federal regulations known as the Common Rule did for ethical review of human subject research.

We have developed such a review process at Stanford. Data scientists evaluate proposed AI tools for bias and clinical utility, and ethicists interview patients and clinicians and developers to learn what they are worried about.

Finally, do not forget about health insurers, for potential harm can result when insurers use algorithms to make medical necessity determinations, as recent investigations of Medicare Advantage plans have shown. In theory, human reviewers are making the final calls. In reality, they may have little discretion to overrule the algorithms.

CMS's final rule addresses this by allowing algorithmic use by Medicare Advantage Plans but requiring them to account for individual circumstances and have a medical professional review each determination. But even as clarified this week, the final rule leaves open important questions about what it means to merely use algorithms as opposed to letting them drive coverage decisions, or to account for individual circumstances, or to have meaningful human review.

So, in summary, to support health-care organizations and insurers, Congress should require that health-care organizations have processes for determining whether planned uses of AI tools meet certain standards; fund a network of AI assurance labs to develop standards and provide expertise and infrastructure for evaluation; require AI developers to disclose necessary information for evaluations; work with CMS on future guidance for health plans; and ensure that Federal agencies have clear authority to require regulated entities to implement these standards. Clarity and specificity are very important here and will be insisted upon by the courts.

Thank you, and I welcome your questions.

[The prepared statement of Dr. Mello appears in the appendix.]

The CHAIRMAN. Thanks very much, and we will have questions in a moment. And I am interested also in getting some of your work connected with what is going on in my State at Oregon Health Sciences University. So, thank you very much.

Dr. Obermeyer?

STATEMENT OF ZIAD OBERMEYER, M.D., ASSOCIATE PROFESSOR AND BLUE CROSS OF CALIFORNIA DISTINGUISHED PROFESSOR, UNIVERSITY OF CALIFORNIA, BERKELEY, BERKELEY, CA

Dr. OBERMEYER. Thank you for this opportunity to address the committee. I am a professor at Berkeley, but my research on AI is inspired by my clinical practice as an emergency physician. I have worked in academic hospitals in Boston and now at a small hospital on the Navajo Nation, and I have seen firsthand how medical innovations can save lives.

But over my 10 years of practice, I have also made a lot of mistakes. Medicine is a hard job, and I wonder, Senator, if you face similar problems in your work. You have to process enormous amounts of information—

The CHAIRMAN. You think? [Laughter.]

Dr. OBERMEYER. You have to process enormous amounts of information. The information is often uncertain and imperfect, and mistakes have enormous consequences. The fact that I have made a lot of mistakes myself is actually what makes me so optimistic about artificial intelligence.

When we talk about AI, it is often very abstract. I want to give a very concrete example of how it can help patients, by telling you about sudden cardiac death. So that is exactly what it sounds like. People just drop dead from fatal arrhythmias, and that kills about 300,000 Americans every year.

The scale of that number is just incomprehensible. Each of those 300,000 is a friend, a family member, a loved one who just disappears, and what makes it even more tragic is, we have a cure. Those deaths are preventable. If we knew those people were at high risk before they died, we would implant a defibrillator to shock their hearts back into a normal rhythm.

So, it is not just that we miss 300,000 opportunities every year to save those lives; when we do place defibrillators, we often place them in the wrong people. About a third of defibrillators end up in people who do not go on to have a fatal arrhythmia.

So, we waste millions of dollars, we expose patients to health risks, and we get no benefit, all because we do not know who is at high risk and who is not. This is a problem that AI can help with. The work that I am doing today that I am most excited about uses patient's electrocardiograms, the wave form itself, to predict sudden cardiac death.

We do a lot better than current methods, which means that one day we can do better in getting defibrillators into the right people, take some of those wasted defibrillators away from people we put them in who are low risk and do not benefit, and give them to some of the people who are at high risk that doctors currently miss.

In health care, it is uncommon that we get a chance to both save lives and reduce costs. Normally we have to pick one or the other, and that is why I think AI is going to be so transformative for our health-care system.

The use cases go well beyond sudden cardiac death, extending to cancer metastasis prediction, Alzheimer's disease. AI is helping to discover new antibiotics, and it can even help identify social vulnerability. We can train an algorithm to spot subtle signals in X-

rays of injuries that can help emergency doctors like me find patients who are the victims of violence when they come through the ER, which is something that we struggle to do today.

Despite all of my optimism, I also worry that AI may end up doing more harm than good if we do not act now. In past work, as you mentioned, Senator, my colleagues and I found large-scale racial bias in the family of algorithms that are affecting health-care decisions for up to 150 million Americans every year.

These algorithms should have been a great use case for AI. They were meant to flag patients who are at high risk of future health problems, so that we can help them with their health needs today. But unfortunately, a design choice in building those algorithms made them biased.

They predicted a patient's health-care costs instead of their health-care needs. Now, costs and needs are very different. Underserved patients, which includes Black patients but also extends to rural patients and any underserved patients, have less money spent on them by our health-care system today because of barriers to access and because of discrimination.

And that means that the AI saw that fact clearly. It predicted the cost accurately, but instead of undoing that in the quality, it reinforced it and enshrined it in policy. Senator, you and Senator Booker sent strongly worded letters to major health insurance companies in the wake of that study, and I believe those had a great impact.

But unfortunately, some of those biased algorithms that we studied are still in use today, and some of their problems surfaced in an investigation of care denial for use of AI in insurance that resulted in harm to vulnerable patients.

I think those examples highlight the need for oversight, and I think there are specific things that this committee and Congress in general can do for the programs under its jurisdiction to help.

First, toward your goal of transparency, I think specificity is incredibly important. We should know exactly what an algorithm is predicting. If it is predicting cost, the developers should not be able to say that it is predicting risk or needs or something else.

Second, accountability. We need to be measuring performance, and especially performance in protected groups under the law in new, independent data sets that the algorithm has never seen, that are diverse enough to reflect the majority of the American population, not just the ivory tower. This is a basic part of good machine learning practice, and we should not take algorithm developers' word that it is performing correctly.

Third, I think that government programs should be willing to pay for AI that generates value and should price those services according to the basic principles of health economics. We do not need to settle for the often poor-quality products that are put in front of us by developers today. We can shape that market, thanks to the purchasing power of those programs.

Thank you very much again for this opportunity.

[The prepared statement of Dr. Obermeyer appears in the appendix.]

The CHAIRMAN. Thank you, Doctor, and I heard you say—and correct me if I did not hear it right. I heard you say there are algorithms in use today that are promoting bias; is that correct?

Dr. OBERMEYER. I believe that is true.

The CHAIRMAN. So, can you give me a ballpark on how many algorithms are being used today that promote bias in health care?

Dr. OBERMEYER. When we did our work in 2019, there was a family of algorithms that included the one that we studied but was not limited to that, that were being used not just by private companies, but also by government programs and academic research groups, that were all predicting costs but being used to make decisions about health.

The estimate at that point was that this was about 150 million patients every year whose decisions about extra help with their health were being affected. Unfortunately, if you look at all of those companies, none of them have publicly disavowed the use of these kinds of algorithms and have clarified that they are no longer using them.

The CHAIRMAN. Well, what was their response to your essentially blowing the whistle? And I apologize to all our guests. I just wanted to kind of freeze-frame this question, because it is so important, since it reflects current bias. So, these were big operations in health care.

These were not like mom-and-pop shops, and you are saying 150 million people were involved then, in 2019, showing bias in health care. You brought it to their attention (a), and (b), they did not seem to do anything about it, according to you. What was their response? Not interested? Who cares?

Dr. OBERMEYER. On one level, their response was very positive. So, in the wake of our study, I actually worked on a pro bono basis with the technical teams at the company whose algorithm we studied, and we rebuilt that algorithm and the same data sets to predict a different outcome, in a way that made that algorithm much less biased.

But since then, it is not totally clear what has happened and whether the current algorithm—or the algorithm that we developed—or the original version is still being sold and marketed for the original purposes.

The CHAIRMAN. So, would it be fair to say that as of today—because you found 150 million patients were being subjected to bias—you have one aspect of this corrected? But from the seat of your pants, and this is not—I just want to make sure my colleagues get a sense of the proportions here.

Your assessment would be probably 100 million patients, even if you take it down some for that one example—100 million patients today are being discriminated against in terms of bias in algorithms. Would that be fair?

Dr. OBERMEYER. I believe that those cost-predicting algorithms are still being used, and I have seen no evidence that they have been taken out of use or are no longer being marketed for the same purpose that they were originally being used for.

The CHAIRMAN. And Dr. Sendak wants to say something, and then we are going to get Dr. Baicker back in the game here. Doctor? I apologize to all of you. I just thought that it was so important

to get a read on the number of people being subjected to algorithm bias.

I say to my two colleagues, thank you very much, both of you, for coming, and I know there has been a lot going on this morning. We just heard from one of the leaders in the field that it was his judgment that 100 million patients are still being victims, as of now, to algorithmic bias.

Dr. Sendak, and then we will move on to Dr. Baicker and my colleagues at the end of this discussion.

Dr. SENDAK. So I just want to put a pin in a point where my colleague Dr. Obermeyer led some of the initial work showing racial bias in algorithms. So, over the last 6 months we have been interacting with a disability rights group.

Obviously, disability status is a protected class, legally, that is named in the Office of Civil Rights final rule or proposed rule. And in our work with them, there are almost no empirical studies looking at disability bias in algorithms that are currently used. So, I mean when you ask about bias, I know that we often think about racial bias. But the empirical evidence to look at other vulnerable groups is significantly lacking.

The CHAIRMAN. Okay.

Dr. Baicker?

**STATEMENT OF KATHERINE BAICKER, Ph.D., PROVOST,
UNIVERSITY OF CHICAGO, CHICAGO, IL**

Dr. BAICKER. Thank you, Senator Wyden. I am honored to be here and have the chance to talk with you all about this very important topic. I am provost of the University of Chicago, and engaged with a number of different health-care organizations, but of course I am speaking only for myself.

I wanted to elevate two themes that we have heard about from our colleagues, and that I know are of vital importance to ensuring that all Americans have the opportunity to benefit from AI algorithms.

The first is that these algorithms can increase the quality of patient care by targeting resources to the patients who need them most. You heard about the case of sudden cardiac death. There are many other examples across the health-care continuum of places where we overuse care in patients where it is not to their benefit and increases cost and decreases accessibility. And we underuse the same type of care for other patients, who go without these vital services that would improve their health. AI algorithms can help us predict who is most likely to benefit from care.

There is enormous risk of bias, as you have heard about, but there is also risk of bias, inadvertently, from physicians themselves. I would be concerned that every patient who sees any health-care provider is being treated in a way that is consistent with that provider's experience with individual patients, with things that have happened recently in that provider's patient panel.

Adding guidance from AI algorithms can help undo the bias that any individual is going to experience. And putting the two together can improve the quality of care and make innovative care more affordable, because even very expensive health care is well worth it

when it improves a patient's quality and length of life, but becomes unaffordable when it is deployed in patients for whom there is really minimal benefit. That raises the cost of insurance and makes health care less affordable for individuals, as well as for Federal programs like Medicare, and Federal and State programs like Medicaid.

Now, if we are going to pay for this kind of care to make sure that the best combination of human and data is available, those algorithms need to be tested rigorously with a broad panel of patients in appropriate settings.

As my colleague noted, the application of these in different settings may have very different results. What looked like it worked well in a small, homogeneous patient panel may have very different effects for different patients. The way in which algorithms are implemented is going to affect how usable they are for clinicians, and therefore how beneficial they are to patients.

So we have a strong interest in rigorously testing these algorithms, the same way we would for any other kind of medical innovation. Beyond individual patient quality and the value of care though, I think there is even greater opportunity to reform the way we deliver health care system-wide.

Right now, there is an under-incentive to invest in care that would improve a patients' well-being over decades. This is particularly problematic for Medicare. Individual insurers or the employers who purchase plans for privately insured employees may not see the return to investing in care that is going to avert a heart attack 20 or 30 years later, when patients are on Medicare.

There is an incentive to invest in care that patients can appreciate the value of in the near term, but often the long-term benefits of the care may be hard to discern, especially when patients are healthy, before they are sick and need the quality of care that the insurance is meant to provide.

So there is a strong public policy role to promote investment in the kind of care that will improve health in the long run, and AI offers us a tool to better capture that and therefore better reward it and incentivize it. Just as risk adjustment now provides a mechanism to deter insurers from selectively enrolling only healthy people, having that kind of population-level long-run risk adjustment can help provide the resources needed to invest in people's long-term health, which is of course first and foremost to their benefit, but also to the benefit of Medicare, from which those patients will eventually be receiving care.

There is also a huge return to public investment in making sure that the data architecture is available to draw in data across silos. The biggest return to these AI algorithms is when you can bring together data from multiple insurers across health and nonhealth-care settings, to really figure out the care that patients need.

That return is not realized by any individual data-gathering insurer or employer, or even health-care researcher. There is a role for Federal policy in saying there is an enormous opportunity to do better for patients and provide higher-value care if we can bring all this information together.

But that also then poses risks to patient privacy, confidentiality, and the risk of discrimination, and so those massive data sets need

to be guarded by much more specific and transparent Federal regulation, to make sure that the data is used to the benefit of all patients, and that is not necessarily the world that we are in now.

So, thank you very much for the opportunity to speak with you, and I look forward to answering any questions that I can.

[The prepared statement of Dr. Baicker appears in the appendix.]

The CHAIRMAN. Thank you, Dr. Baicker. We look forward to questions from our colleagues. Thank you all for coming.

It is obviously a busy day, and I want to say to our guests, normally the chair asks the questions first, but in deference to my colleagues—because I have been at it basically for an hour or so—I am going to let all of the Senators ask their questions first, and next in order of appearance will be Senator Menendez.

Senator MENENDEZ. Thank you, Mr. Chairman.

Dr. Obermeyer, there is a growing concern that algorithms may produce racial and gender disparities via the people building them or through the data used to train them. For example, health systems often rely on commercial prediction algorithms to identify and help patients with complex health needs.

A study conducted by you and your team found evidence of racial bias in a widely used algorithm. Because this algorithm used ineffective proxies and falsely concluded that Black patients were healthier than equally sick White patients, Black patients were significantly less likely to be identified for extra care. As Congress considers appropriate payment and coverage policies for AI, we need to ensure that AI is not building upon biases in health research, and compounding health equity issues.

What steps can policymakers take to ensure that AI can be used to improve health outcomes for underserved and underrepresented populations, rather than build on the current health disparities?

Dr. OBERMEYER. Thank you for asking, Senator. I think the two principles that I would say are important, that I believe Congress should enforce for algorithms that are being used on patients are, number one, transparency in the form of specificity.

We should know exactly what algorithms are predicting. So, if an algorithm is predicting cost, it should say that very clearly. And I think that will be part of educating the market of health systems and others who use those algorithms, that if this is just predicting how much someone is going to cost, I probably should not use that to figure out how much health care somebody needs, because there are some people who are low-cost who face barriers to access and discrimination that has reduced that cost, who actually need more care than their costs would indicate.

So transparency, in the form of clearly labeling what an algorithm outputs, is important.

The second is that we need more accountability in the form of evaluating those algorithms in new data sets and by third parties, so that we do not have to take an algorithm developer's word that the algorithm is working well and equitably across groups.

Senator MENENDEZ. Thank you. Let me just stay with you for a moment. Lack of diversity in clinical trials—something I have been working on for quite some time here—creates gaps in our understanding of disease prevention and treatment across populations.

Legislation like my bipartisan DIVERSE Trials Act seeks to improve access to and diversity in clinical drug and treatment trials, but more needs to be done. Now recently, some stakeholders have been using AI to increase diversity in clinical trials by pinpointing community centers where patients with certain cancers might seek treatment.

That information helped lift the Black enrollment rate in five ongoing studies from roughly 4.8 percent to about 10 percent. How can we support stakeholders looking to invest in AI to improve diversity and ensure medicines are representative of a broader population?

Dr. OBERMEYER. Thank you, Senator, for pointing out that AI can actually have huge returns in increasing diversity and equity in the science underlying medicine. I think finding the types of patients who do not typically get enrolled in trials today is a very important part of that, and I think, as Professor Baicker mentioned, linking data sets together to find those patients, to use AI to identify patients the doctors are currently missing for trial enrollment, has great promise for improving the diversity of the science on which clinical medicine is based.

Senator MENENDEZ. Well, thank you for that. We look forward to working with you and some of your colleagues.

Mr. Shen, according to a recent study, there are fewer than 20 AI medical services for which CMS reimburses today. The payment for those services—which include AI systems for cardiology, ophthalmology, and radiology—is inconsistent and varies widely.

In 2023, most AI medical services had fewer than 1,000 claims billed nationally. In short, the AI medical services that we have online today are just beginning to be widely accessed by patients. Would more consistent Medicare payments make these services more accessible for patients, and are there any other holes in the way CMS is currently reimbursing for Algorithm-Based Health-care Services?

Mr. SHEN. Yes; thank you for the question, Senator. So, while CMS has recognized the value and the complex nature of FDA-cleared AI solutions, the agency's reimbursement decisions have not been uniform or consistent in terms of ensuring appropriate levels of payment for those products.

We support a solution that ensures a predictable and consistent approach by CMS, an approach that recognizes the cost of AI used in the clinical patient care and that reimburses with a temporary and separate payment for the distinct service clinical AI tools like Algorithm-Based Health-care Services provide, which otherwise would be unavailable based on the qualitative and quantitative AI analysis that those solutions provide.

More specifically, we advocate for CMS to implement a consistent and reliable payment policy for technologies and services that are first, FDA cleared; second, a covered benefit under Medicare; and third, that the service must provide a clinical output that supports the physician's decision-making.

Furthermore, we encourage Congress to encourage CMS to formalize its existing Software as a Service policy, which would allow for separate and distinct payment for Algorithm-Based Health-care

Services with at least 5 years of consistent payment, while given a new technology payment assignment.

We believe all of this will allow for better adoption of Algorithm-Based Health-care Services to help CMS collect more data to evaluate the overall value of AI to patients. We also believe that more data will demonstrate AI's ability to increase access to care and improve patient outcomes.

Senator MENENDEZ. Thank you, Mr. Chairman.

The CHAIRMAN. I thank my colleague.

Senator Cortez Masto?

Senator CORTEZ MASTO. Thank you. I want to thank the chairman, ranking member, and the panelists. This is a very timely and important discussion. If you have not been to CES, the Consumer Electronics Show in Nevada, I would welcome you there. Cutting-edge technology in this space that we are talking about right now in the use of AI, the impact it has to our patients, but also the health-care industry and access—

But this is a concern for me: how do we overlay necessary regulation when it comes to AI, necessary regulation and guard rails for data privacy? Because I think that is all part of it when we are having this discussion. But one of the things that also comes up—and, Dr. Mello, I am going to ask you if you would address this, and any of the other panelists are welcome.

I have been reviewing recent investigative news coverage of the use of AI and algorithms in limiting access to skilled nursing services. We have heard this happening in other post-acute settings as well, like home health and rehab care. I share concerns expressed here today about Medicare Advantage plans' use of AI tools in prior authorization decisions.

So, Dr. Mello, I appreciate you highlighting these issues with AI and health insurers in your testimony. As you note, CMS finally recently finalized a rule requiring Medicare Advantage plans to have a medical professional review in AI-assisted decisions.

So my question to you—and I have two of them. And one is, in your view, is human review an adequate guard rail to ensuring plans are individualizing coverage decisions, and does the CMS rule provide enough oversight of the use of AI tools in coverage decisions? Those two questions. Thank you.

Dr. MELLO. Thank you for that important question. I think there are few alternatives to human review. So that is where we ought to focus. The question that interests me is, what does meaningful human review look like?

Senator CORTEZ MASTO. Right.

Dr. MELLO. As you may have heard, there was another insurer that used a non-AI-based algorithm to deny care. That did have human review, but the human review, on average, took 1.2 seconds. And the CMS final rule currently does not include the level of specificity that would help plans understand what meaningful human review looks like.

Senator CORTEZ MASTO. Right.

Dr. MELLO. In order to enforce incentives to make it meaningful, the second point I would make is that audits by CMS need to look very closely, as I believe they intend to, at denials where algorithms were involved, to require transparency about when algo-

rithms were involved and to really look at the patterns of denials and reversals.

Senator CORTEZ MASTO. Okay, and this goes to, I think, my colleague Senator Menendez, and I agree with this. As we are looking at AI tools and we are trying to address what we see may be discriminatory impacts of the data, our use of AI is only as successful as the data that we are relying on.

If the data already has a bias for discrimination or the like, that is a concern. And for us—and this is why this is an important discussion—what role does Congress have to help regulate that, prevent that from happening? I am going to open this up to the panel to please weigh in.

But let me just—I only have 2 minutes left here. Mr. Shen, let me talk to you about this issue as well. And Siemens—thank you very much—is in Nevada. I have been there. I so appreciate the work that Siemens does.

A recent survey of health system executives found that, while many believe generative AI has the potential to reshape the industry, only 6 percent have established a generative AI strategy. It is surprising to me that we are talking about all of this, but the uptake is so low in the health sector.

Maybe that is not a bad thing, as we are talking about the challenges we are facing right now. But there is also potential, positive, for AI to do really good work in this space. So I guess my question to you is, can we talk a little bit about Siemens' work to collaborate with physicians in developing these AI algorithms, and how has this partnership also impacted that adoption?

And then finally, as we are having this discussion, protecting that data and the algorithms are only as good as they are clean and take out that internal bias. How are you addressing that?

Mr. SHEN. Yes, thank you for the question, Senator. Certainly, I think working with our clinical partners is very critical, not just in the development of our AI algorithms, but also in the adoption of those AI algorithms as well.

So, from our perspective, we work closely with those clinicians to really identify what is the need that they have. What is the challenge that they have from a clinical standpoint that we should try to create that AI algorithm to help them address?

And more importantly then is, what is the patient population that they are representing, so that the data that we use to train that AI algorithm is reflective of that patient population? And then, once we actually apply that AI algorithm into clinical practice, it is important that we facilitate the adoption.

So what does that mean? That also means educating the clinician in terms of what is the intent of that AI algorithm. What was the purpose, as some of my colleagues have talked about here, of that AI algorithm? And then also to provide transparency as to why is that AI algorithm making the clinical decision or recommendation that it is making—so, giving that education to the clinician as to why that is happening.

And finally, I think the other important aspect is that we work very closely with our clinical partners to gather feedback from them, to really understand how do we improve and continue to innovate that AI algorithm going forward? These are all critical steps

where we believe we need to have a key partnership with the clinician.

The CHAIRMAN. We are going to have to move on, Mr. Shen. Thank you.

Senator CORTEZ MASTO. Thank you, Mr. Chairman.

The CHAIRMAN. We have a lot of our colleagues. I thank my colleague.

Next is Senator Johnson.

Senator JOHNSON. Thank you, Mr. Chairman.

I have heard two witnesses mention the need for transparency, which I totally agree with. But I want to just tell a little anecdote about the lack of transparency in government.

CDC/FDA has what they call FAERS and VAERS systems. Prior to Emergency Use Authorization of the COVID vaccine, the FDA/CDC was touting the Vaccine Adverse Event Reporting system. You know, they are going to be watching this. They are going to survey it for safety signals. You lose a couple of days away from work, they are going to have a CDC representative call you and follow up on you. That was total BS.

After the vaccine was rolled out and they did not like what VAERS was reporting, they started denigrating the VAERS system. They also created a standard operating procedure which talked about how they were going to analyze that data, first talking about proportional reporting ratios and then empirical Bayesian analysis.

First, they denied actually conducting those analyses; then later we found out somebody said they actually had done it. So they agreed that, yes, they had actually done it. So I have been, for over a year, trying to get the information from CDC and FDA in terms of their empirical Bayesian analysis of the VAERS data.

Now again, these are agencies that we fund, and we pay for—taxpayers. These are government employees we pay the salaries for. This is a standing operating procedure that they are going to do analysis on the VAERS system. They will not turn it over to me.

So I just mention that in terms of anybody thinking, well, what we need to do is, we need to create AI and government regulations so that we ensure transparency. Nothing could be further from the truth. I am old enough to remember, before we started calling this AI in medicine, we talked about expert systems.

I remember a PBS series one time or a report, and I was really surprised by the doctors they were questioning about expert systems. They were really opposed to them. It did not make sense to me. I thought, you know, why not? I remember Dr. Dean Edell was always talking about, when you hear hoof prints, think horses not zebras.

Well you know, every now and again it is zebras, and an expert system, that kind of data, could potentially give you that information: drug interactions, that type of thing. I think during COVID—again, I have always been supportive of things like expert systems, AI systems having access to anonymized data. I think that is the real key: how do you anonymize this data so that researchers, a bunch of them, can? I think the danger of AI is central control, and the solution—nobody knows exactly where this is going to go—is as

much AI out there, you know, as many systems as possible kind of watching each other, trying to——

The same thing goes with medical protocols. I mean, it makes an awful lot of sense to have basic medical protocols, but you also have to let doctors be doctors. You have to let them utilize the tool of an expert system or an AI, but then use their training and their medical system and be able to practice medicine, as opposed to having it dictated.

I think one of the big problems in medicine today is, we do not have very many independent doctors now. They are all part of the system, and they all have to follow exactly what the system tells them to do. And of course that system does exactly what the Federal health agencies tell them to do, and the Federal health agencies are not transparent.

We have a real problem. We have seen a real corruption of medical science research. It has been a corruption, with big pharma controlling all these things and relying everything on random controlled trials that only they can pay for, completely ignoring the use of other molecules of generic drugs because nobody wants to pay for that.

So again, to me, the real benefit of AI with anonymized information—I think Epic Systems has a system. They are trying to do it, and there are clients that agreed to provide their data, anonymized. They can go then and do research. I mean, I think that is the kind of model we have to work toward.

But I think the main question is, how do we allow that system to flourish, and how do we protect individual doctors, individual researchers, and not have them come under the control and under the thumb of the medical establishment that has been, I would say in many respects, corrupted by big pharma? Whoever wants to take the question. I'll go to Dr. Baicker.

Dr. BAICKER. Well, thank you. I think it is vitally important that there be transparency in how algorithms are deployed and in the data that feeds into them. There are some promising examples, I think, of public-private partnerships there. Where the cutting edge of aggregating and anonymizing data is, I think at this point—although I will turn to my colleagues who have expertise in this as well—is really sitting in the private sector in health systems, in nonprofits, in entities that have some degree of trust in being able to hold data that is anonymized, but merged, so that researchers can explore opportunities to find new cures, to find new pathways for patients, that takes the individual's identity away, but uses the full set of information.

That is vitally important to getting those insights, that assurance that you are pointing to. We need assurance that algorithms are actually delivering the outcomes that we care about versus other outcomes that might be easier to find—and that they are representative of patient pools broadly across the country, not drawing just from one type of patient and applying the insights to other types of patients for whom they might do more harm than good.

Senator JOHNSON. So just with my final time. So now we are talking, with AI, we are talking about observational data, which the Faucis of the world just denigrated, because only random controlled trials——

So I mean, if we are going to really utilize this, we have to bolster, again, the value of observational trials. And we need a lot of people looking at the same data, and then have that data transparent. These trials, they are not making—even peer-reviewed, the reviewers do not get access to the data, and that is just wrong.

The CHAIRMAN. I thank my colleague.

Senator Bennet is next.

Senator BENNET. Thank you, Mr. Chairman. I am grateful to you and the ranking member for holding this hearing and giving us a chance to ask questions. And thank you to the witnesses for being here today.

I wanted to shift this, Dr. Baicker, in a slightly different direction than we have covered as I understand it today, and that is the issue of cost in artificial intelligence. It is hard to believe this, but we spend \$4.5 trillion, which is more than 17 percent of our economy, on health care.

That is twice as much as any other industrialized country in the world, and I think it is worth asking—it is always worth asking what we are buying for that. Our life expectancy at birth is 3 years lower than other developed countries. We have the highest rate of infant and maternal deaths of any industrialized country in the world. We have among the lowest rate of practicing physicians.

This is an incredible statistic, I think. We have among the lowest rates of practicing physicians when standardized for population size, and some of the lowest rates of per capita physician visits in the industrialized world as well. People are told that they should suffer the inconvenience of our system because it is easier for us to see a doctor in the United States, and that actually turns out not to be true. We have the inefficiency of this system that we have, plus nobody can get in to see a doctor.

The same holds for patient hospital stays and the number of hospital beds. We are outperformed by other countries around the world, even if we are spending twice as much as they are on health care. We are spending enormous sums of money, I think, without seeing the results that the American people deserve.

Not only is this bankrupting working families, it is adding billions of dollars to our national debt. Part of this is due to the administrative costs that are spread across our health-care system. I think that is a big piece of that, and it seems to me that the careful deployment of AI tools that prioritize patients might help, and maybe this is an area where we could actually prioritize patients and not profits in the system. They could remove expensive intermediaries and reduce the paperwork burden—and create new efficiencies that could drive down costs.

So, Dr. Baicker, I would ask you, how can AI tools address administrative expenses and reduce the costs we see in our health-care system, while improving patient cost of care?

If there are other folks on the panel who would like to get into this conversation on administrative or any other aspects of the health-care system, please feel free. Thanks, Dr. Baicker.

Dr. BAICKER. Well, thank you so much for the opportunity to talk about this crucial issue you are raising, that the U.S. health-care system is not delivering the quality of care and outcomes for patients that we ought to expect, particularly given how much we are

paying for it. I do think AI offers some opportunity to improve patient outcomes while slowing the growth of health-care spending by being sure that we are not providing unnecessary care that is potentially harmful to patients, and freeing up those resources to do all of the things that we need to for patients' health that we are not delivering consistently now.

I think sometimes there is, in the public discourse, some conflation of high-value care and low-cost care. There is lots of care that is of very high value and very expensive that produces wonderful outcomes, or very high-value and low-expense that produces outcomes in a really cost-effective way.

We want to target resources toward where they produce the most health, that mix of high-cost and low-cost care that is really appropriate for patients. AI can help us do that. But even more, I would like to see us move toward a system of paying for value rather than paying for the inputs into health care, and paying systems that produce better patient outcomes more, rather than just paying based on how many resources they use.

AI can be part of that in that if AI helps a doctor, an insurer, a system, deliver better outcomes for patients at lower cost, that should be an advantage there, and we can pay for bundles of care that produce patient outcomes rather than each of the individual inputs.

Senator BENNET. Is there anybody else who would like to take my last 30 seconds? Dr. Obermeyer?

Dr. OBERMEYER. Just one thing. Through my research and because of family, I spend a lot of time outside of the U.S., and there is actually nowhere I would rather get medical care than in this country.

Now I am lucky, because I am a doctor. I am privileged in lots of ways. I can navigate the system in ways that others cannot. So I think AI has a lot of potential to help everybody, not just people like me, navigate that system.

A lot of the administrative decisions, even something as simple as, you want to see a doctor; when should you see that person? How early? That looks like an administrative decision, but it is actually also a health decision, because it depends on that patient's health needs.

So I think AI will be particularly powerful at that interface of these decisions on the back end of health care, like population health management, scheduling, et cetera, that are actually both administrative decisions and health decisions. AI can be very powerful there.

Senator BENNET. Thank you. Thank you both, and thank you to the panel.

Thank you, Mr. Chairman.

The CHAIRMAN. Thank you, Senator Bennet. I would just say, the first thing Senator Bennet said—and he is always very informed on these issues—is also worth just a quick comment.

He mentioned that we are spending \$4.5 trillion a year on health care. And ever since I was director of the Gray Panthers, what we would do is divide the number of Americans into the collective spending. Today, far different than when I was coming up, \$4.5 trillion divided by 330 million Americans means that what Senator

Bennet was just saying, that if we wanted to do it—and this is not how we are approaching health policy—but if we wanted to do it, we could send every family of four in America a check for \$50,000 and say “good evening, happy to be here”—\$50,000 for a family of four.

So Senator Bennet, as usual, gets us off on the right track. I thank my colleague.

Senator Young, welcome.

Senator YOUNG. Well, thank you, Mr. Chairman. I want to thank you and members of the Finance Committee staff. As the chairman knows, I have been working on artificial intelligence, trying to learn everything I can, working with Senators Schumer, Heinrich, and Rounds.

We have held a number of forums on the topic. We have attempted to educate some of our colleagues on what adoption of AI technologies is going to mean for all various facets of life, from national security to labor market impacts to health care.

And I have learned quite a lot, and the hope has always been that our committees of jurisdiction would sort of grab the bit and run with it, and the Finance Committee is out front here. So, thank you to our witnesses for helping us sort through this issue.

I would agree with Dr. Baicker. You know, I have heard a lot of commentary about the productivity improvements, but also the health and wellness improvements we might realize, and perhaps we will actually be able to bend the proverbial cost curve down in health care. We have tried seemingly everything else in Washington, so maybe innovators and entrepreneurs and investors can help us get there.

So I will begin with some questions for Mr. Shen, please. Mr. Shen, on issues of workforce, the Federal workforce, are there existing gaps in the CMS workforce and expertise that may create choke points to the access or the use of AI within the health-care context?

Mr. SHEN. Thank you for the question, Senator. Certainly, I think CMS has recognized the value and the complex nature of FDA-cleared AI solutions, but the agency’s reimbursement decisions have not been uniform and consistent to ensure appropriate levels of payment for these products.

I think what we are advocating for here is to ensure some sort of predictability and consistency as an approach from CMS, an approach that recognizes really the cost of AI and how it is benefiting the clinical patient here, and these different clinical AI tools that are providing both quantitative and qualitative information for that physician to better help the patient.

We see an opportunity here to actually create a temporary and separate payment for the distinct service that these clinical AI tools are providing, that would otherwise be unavailable for the physician. And specifically what this committee can do is really to advocate CMS to implement this consistent and reliable payment policy for these FDA-cleared solutions that we have that are affecting clinical AI.

Senator YOUNG. Thank you. Thanks for the actionable recommendation.

How can Congress—as it relates to maintaining our competitiveness, how can Congress support continued innovation, responsible use of AI, as well as continuing to assess the future risk? Relatedly, how can we partner with industry to advance these goals?

Mr. SHEN. Yes, that is a great question, Senator, and thank you for that. So what we see as a real key component for innovation is adoption, and we believe that the adoption of these new and emerging technologies like artificial intelligence will fuel innovation, innovation here in the U.S., innovation by technology companies like ourselves, to be able to continue to push the forefront around artificial intelligence.

So a key component again here is really driving adoption. So, enabling providers to be able to embrace this technology without any uncertainty in terms of making the investment in this technology, and by encouraging adoption, that is going to fuel innovation going forward.

Senator YOUNG. Thank you.

One elementary sort of piece of information about artificial intelligence which I learned early on is to think about the technology. There are three elements. There are computers, that is processing power; there are those who actually train models, so talent; and then there is data.

You need a lot of data oftentimes to train these models. You of course need appropriate and clean data. Innovation is rapidly changing what is possible, and data sharing can help leverage advancement of these AI tools. How do we balance necessary data sharing with protection of the individual patient data, and protection of large data sets?

Mr. SHEN. Yes; another great question, Senator. I think it is very important that the data sets that are being utilized to train these AI algorithms are inclusive and reflective of the patient population that the AI is going to be applied toward, and that is a key component here—so, making sure that the data that is being utilized here is focused on the patient population that it is serving.

Senator YOUNG. That is great. Thank you.

The CHAIRMAN. I thank my colleague. We have four Senators here, and obviously noon is really the deadline. So, we can get everybody in, and we will. Let's see—Senator Carper is next, followed by, at this point, Senator Cardin, Senator Blackburn, and Senator Thune. And if other Senators come, we will figure out how to deal with that.

Okay; Senator Carper?

Senator CARPER. Thanks Mr. Chairman. Good morning, everyone. Thank you for spending your time with us today. Thank you for your work in what is really an increasingly important issue. We have gotten our heads around it. There is a lot of potential here, also some real pitfalls if we are not careful.

So, I want to ask some questions. One of them deals with the workforce and the responsible deployment of our workforce. Strengthening the health-care workforce is an issue of key concern to me, and I know it is to Democrat and Republican colleagues.

I consistently hear from providers back in Delaware and other places too, and from health-care systems, that recruiting and re-

taining a health-care workforce is the number one issue that our health-care system faces.

In fact, I go home almost every night to Delaware. When we are not in session, I cover my State. I call them customer calls, sometimes to health-care places like the hospitals and all, but all kinds of places, businesses large and small.

And I ask three questions when I do these customer calls. I say, how are you doing, the business? How are we doing, the congressional delegation, Congress and State Government and so forth? And what can we do to help you? What we hear most, almost from everyone is, we need people who will come to work, people who are trained, people who are trainable and can help us do the work that we are doing.

But I think with the growing use of AI in health care, it is important that this technology be deployed in, I will just say a responsible way that bolsters our workforce, while improving the patient experiences and hopefully outcomes.

We have heard that AI has been referred to as a copilot. I am a naval flight officer, a retired Navy Captain. I spent a lot of time in airplanes, but we hear AI referred to as a copilot. And as we learn of the ways that AI can support our health-care system, we must ensure that our health-care providers remain the pilot—not the copilot, but the pilot of these incredibly powerful tools.

My staff, in fact my colleagues, oftentimes hear me say “find out what works; do more of that.” Find out what works; do more of that. We have seen health-care leaders successfully deploy AI tools, including the health-care systems in Delaware and other places, in a way that has improved patient outcomes and mitigated provider burden.

Dr. Sendak, what best practices should we consider on this committee to ensure responsible deployment of provider-led AI?

Dr. SENDAK. Thank you for the question. So I will comment on kind of two extremes that I have seen. In my testimony, I refer to what I have been able to be a part of at Duke, which is really an exemplary organization investing in the technical capabilities. We have engineers, data scientists, data engineers. We manage projects. We are funded by the health system.

That works. That works to develop, to implement, to monitor AI solutions. That is not scalable. That is not an investment that every organization can or should make. So, when we started Health AI Partnership back in 2021, we were focused on education.

We thought, let us create the best content that we think we can build, put it out there for free, and build relationships with organizations so that they are aware of it and they can adopt it. So, our original scope was curriculum development.

I think we put out some really great content. We interviewed close to 90 people across health-care organizations. We identified eight key decision points along the AI product life cycle. We put out guides to do 31 discrete tasks.

What I quickly realized is, education is only relevant if you have people who can take it, people who have time, people who have the foundational knowledge to go through that training and then bring that back to their organizations. And so, while the content is great, we found that we needed to go deeper.

That is why one of my main recommendations in my testimony that I will go back to is technical assistance. So, a program that we are going to be launching soon we are calling the Practice Network. That is based on other models across the country of kind of hub-and-spoke support, where you have a Center of Excellence that provides specialized support to low-resource settings. I referenced Regional Extension Centers.

So, I think if you are asking to find out what works, do more of that, we have seen exemplary examples of this. I will name Project Echo out of New Mexico that now is used for many different medical conditions. We started to do that type of hub-and-spoke model for AI, and it cannot just be content.

Senator CARPER. All right. I think my time has expired. I have some more questions, but I will submit those for the record.

The CHAIRMAN. I thank my colleague.

Senator Cardin is next.

Senator CARDIN. Thank you, Mr. Chairman. I also want to thank you for holding this hearing, and I thank the leadership for really getting us all engaged in AI.

I thought Senator Young's comment about all of us trying to understand AI and then having our committees try to take appropriate actions in order to have the guard rails that are necessary as we deploy AI—that seems to be the general consensus here. We are all struggling as to how to implement that.

So, Dr. Obermeyer, let me start with you first, because I was listening to your response to Senator Menendez's points on the equities, and there are many issues involved in AI and medical care.

We hear about designing AI to protect privacy. We hear about designing AI in order to protect security. But I do not hear much about designing AI to protect equity for the underserved communities. So, if we are interested in making sure that we really use AI in a fair manner for all populations, how do we incorporate that in the design?

Dr. OBERMEYER. That is a really great question, and I would say that most of my work has actually tried to answer it. I think one observation is, AI is just data. It learns from the data that you give it, and we always feed data into AI that reflects the past, the way our medical system looks today and the way it has looked. That is how AI learns.

Unfortunately, if AI learns to basically replicate our current system, it is going to replicate all of the inequalities in our current system. And so, if AI predicts who doctors should see by looking at who doctors have seen in the past, that is going to incorporate all of the biases that reflect the barriers that people face getting into the health-care system.

So, when we design AI, when I think about designing AI in my research and my applied work, I try to think about how we get AI to reflect the patient's health, not a doctor's opinion. So, if we can actually design AI to predict objective biomedical outcomes that a patient is going to experience and feed those predictions back to a doctor in a way that helps the doctor say this person is at low risk and does not need a defibrillator to be implanted in their heart, this patient is at high risk and does need a defibrillator to be implanted on the basis of a lot of data that that doctor cannot process,

that is a lot better than asking an algorithm to predict who gets defibrillators today, and let's just keep doing that.

So, designing for equity means designing algorithms that look at patients and represent their health conditions in an accurate way and see the individuality of the patient, not just replicating all of the biased and flawed decisions that are currently being made in our health-care system.

Senator CARDIN. And who sets those guidelines or parameters to ensure that the design systems incorporate that?

Dr. OBERMEYER. I think one of the challenges is that every AI is a snowflake. There are so many different uses, and that is good. We want a lot of innovation. But the problem is, it means we cannot set one guideline for all of AI.

So that is why we need, first, a lot of transparency, so that everybody can look at the AI and see, okay, what is it predicting? Is it predicting a good thing or a bad thing, and let me make that judgment. The second thing we need is for that algorithm to be subjected to rigorous evaluation, and so we need that algorithm to be tested in data that it has never seen before.

We need to see how it is performing not just overall, but in Black versus White patients, urban versus rural, all of these axes of inequality that we currently have in our health system, to make sure that the algorithm is not just performing well overall, but it is performing well for everyone.

Senator CARDIN. Dr. Sendak, if I could follow up on a similar concern on equity, and that is that we have a digital divide. We have challenges with accessibility for underserved communities generally with health care, but also our safety net providers do not necessarily have the same capacity.

So how do we build into the system that AI will be really accessible, particularly in underserved communities?

Dr. SENDAK. And I want to just build off Dr. Obermeyer's comments. I want to emphasize that even when we went out, and in our interviews we asked 10 organizations how do you assess for equity in algorithms, we got different responses. So, we have tried to start to do that consensus-building, but it will take investment to build those capabilities.

And so to your point, low-resource settings—and I mentioned this in my testimony—today, they are not on the AI adoption highway. Guard rails do not help them, because they are struggling to operate efficiently at baseline.

So that is where I want to keep pointing to prior Federal programs where we made infrastructure investments in procurement of new digital platforms to enable care to be delivered with EHRs. We invested in technical assistance programs.

So that is how we are going to reach those communities—and just going back to getting it out of the ivory tower in those settings.

Senator CARDIN. Thank you.

Thank you, Mr. Chairman.

The CHAIRMAN. I thank my colleague.

Next is Senator Blackburn.

Senator BLACKBURN. Thank you, Mr. Chairman. Thank you to each of you for being here.

I had a great roundtable last Friday with health-care innovators in Nashville, and of course we've got great work that is being done in health IT there, a lot of AI that is being done. HCA, one of our hospital corporations, has really deployed so many AI tools to help with clinical decision support, nurse handoff, documentation, things of that nature.

One thing that came up was CMS and the reimbursement framework. And, Mr. Shen, I do want to come to you on this, and I know it has been brought up before.

There are deficiencies, there are discrepancies. You have coding that calls for this, this, and this, and then over here you have this new technology. So this came up at the roundtable. So, talk to me a little bit—let us drill down a bit more on these deficiencies and these gaps, and how can CMS move forward and fix that, so that patients have the access to the best possible—

Mr. SHEN. Yes, I really appreciate that question, Senator. I think you have identified the issue spot-on here. With CMS, we are seeing inconsistencies and an inconsistent message in terms of how do we appropriately reimburse for these AI technologies.

From a practical standpoint, what we are advocating for here is kind of a predictable and consistent approach that is really focused on clinical AI solutions. So these are really Algorithm-Based Health-care Services that are providing both quantitative and qualitative clinical information to the physician, so that he or she can be—

Senator BLACKBURN. Should CMS be more prescriptive in how this gets handed back to them for reimbursement?

Mr. SHEN. Yes. I think that what is plausible here is that there are pathways that are already in place. There is infrastructure already in place within CMS—

Senator BLACKBURN. So it is our job here at the Finance Committee to tighten this up?

Mr. SHEN. To really push them, yes, in terms of consistency. I am really advocating now for reimbursement for, again, these FDA-cleared clinical solutions that are already covered under Medicare.

Senator BLACKBURN. Okay; all right.

Then, Dr. Mello, same type thing on regulation. When you look at the regulatory frameworks and the balance that needs to be there to establish those standards, talk to me about how this keeps up with that evolving landscape?

Dr. MELLO. So, I think what your question raises, Senator, is, how do we get this balance between innovation and patient protection right? And I think the critical step that CMS and other agencies—and Congress—could take at this stage is to convene experts in the field to begin to establish consensus-based standards about how algorithms need to perform, and about how the people who use them need to plan for their deployment.

So, CMS could do this, for example, as part of Medicare conditions of participation. I think that would be a terrific way to get the attention of health-care organizations that are using AI tools but not planning for them in the way that they should, and not monitoring them.

Senator BLACKBURN. Dr. Baicker, you talked some in your testimony about AI's potential to overhaul health care payment sys-

tems, and as we look at value-based payment models, we are there with you on that. It does not matter if it's blockchain, if it is AI, but achieving that potential—

So how do you envision policymakers really harnessing AI effectively to refine these payment systems?

Dr. BAICKER. Well, I think there are two mechanisms that might be low-hanging fruit in that way. One is that we do not really pay for improving, avoiding future health risk.

Right now, if an insurer does a great job at reducing the likelihood of future heart attacks, there is really very little payment associated with that. So there is a real barrier to investing in that kind of care. AI can generate markers of future health risk that we do not currently have—

Senator BLACKBURN. The predictive diagnoses?

Dr. BAICKER [continuing]. Then we can pay for it, once we have a better marker. So that is one mechanism. Another is in sort of smoothing the way that we deliver care now. We can actually reimburse the quality of care that patients are getting today in a much more nuanced way, drawing on the data that we can now process with AI.

All of these data sets have existed before, but it was very hard to boil them down into something you can use.

Senator BLACKBURN. My time has expired. Thank you.

The CHAIRMAN. I thank my colleague.

Next is Senator Whitehouse.

Senator WHITEHOUSE. Thanks, Mr. Chairman.

One of the unfortunate things about these hearings is we have a really terrific panel of experts on a really complex subject, and 5 minutes. So the questions that I will ask, unless somebody has an immediate response, I would like to have each of you take for the record, so you have a chance to deliberate a bit, think about it, and get back to me.

I represent Rhode Island. Rhode Island has several health-care enterprises that we are very proud of. One is what is perhaps the best all-payer claims data base in the country, and with that array of data information there is obviously a role for AI in trying to plow through it and seek out either anomalies or associations that might not be immediately apparent.

I would be interested in—question one: is there any advice that you would have with regard to maximizing the utility of a robust all-payer claims data base? So that would be question one.

Question two is, are you working in the development of AI with some of the medical specialty groups—orthopedics, cardiologists—and do you think they are a useful place for either benchmarking or approval or somehow accrediting best AI practices? And if they are, do you have any good examples of a medical specialty association that is being particularly forward and helpful at looking for the best uses of AI within the specialty? So that would be question two.

The third question has to do with ACOs. Rhode Island has two unusually good Accountable Care Organizations. One is Integra, and the other is Coastal Medical. They are national leaders in terms of quality plus savings, and they have done really, really, really well.

I would be very interested in your thoughts on ways in which artificial intelligence could be used to support the ACO program. To me, the ACO program is one of the most exciting things happening in medicine right now, just the effect that it has on incentives, and the way it opens up a medical practice to see the value in nontraditional interventions that will actually help the payment and save money but would never have been reimbursable under a traditional pay, fee-for-service system.

So it is in that area, I think, that there is an enormous opportunity. ACOs can be very big, but they can also be pretty small. And I think the best ones, at least in my experience, have been provider groups that have not tried to pull together too many different interests and struggle through the conflicting interests, but just to be a good ACO provider group providing primary care.

So, if you could focus your responses a little bit on that type of ACO, a primary care provider ACO, and what I should be looking for or encouraging or trying to support, so that we can continue to advance the ACO model as we try to redesign health care.

I have 2 minutes and 30 seconds left, so if anybody wants to jump in and use that 2 minutes and 30 seconds on any one of these. Okay, go ahead, but these are questions for all of you, and I would be really grateful if you would take a moment when this is over and send a written response to me. Go ahead.

Dr. SENDAK. So I will bucket the last two, and I will speak now from the perspective of Duke. So many of the initial use cases of AI that we were involved in, in terms of projects that implemented AI, were within the ACO population.

And to your point, that was a unique environment where prevention of complications that you could build algorithms to predict actually generated——

Senator WHITEHOUSE. Revenue.

Dr. SENDAK. Revenue.

Senator WHITEHOUSE. Yes. It is a beautiful thing——

Dr. SENDAK [continuing]. And facilitated the investment in the technical infrastructure, in the clinical expertise required to prevent those complications.

So, to loop in your second question about specialties, every year our team at Duke DIHI runs an innovation competition where we ask people, okay, bring us your problems and we will help you identify solutions that we can implement.

So, what is very, very common is that specialists come to us. It could be an ophthalmologist with diabetic retinopathy. It could be a vascular surgeon with peripheral artery disease. It has been a nephrologist for chronic kidney disease.

They come to us, and they say, hey, I see all of these complications from chronic illness. If the primary care doc did x, y, and z, we could have prevented this. So everyone always points upstream to that primary care doc. So what we have done across multiple chronic conditions is build algorithms that prevent that complication that the specialist has to take care of.

We developed a workflow called Population Rounding, where that specialist actually reviews high-risk cases that are identified by algorithms, and then they collaborate with primary care docs to send

them recommendations that the primary care docs intervene on, and generate revenue for the ACO.

The ACO pays for that specialist's time to do the case reviews, and they pay for the technology. That has been sustained now for 8-plus years for multiple conditions at Duke.

The CHAIRMAN. The time of the gentleman has expired.

Senator WHITEHOUSE. Well, that is really interesting, because one of the problems that we face with the ACOs is that they are heavily leveraged. Coastal at one point told me that the primary care doctors at Coastal Medical controlled, I think he said, 13 percent of the care of patients.

But the ACO calculation was responsible for 100 percent, with the rest being pharmaceuticals, specialists, hospitals, and other things. So the more we can integrate specialists and primary care into that ACO model so it is little bit more like a general contractor thing—

I would be really interested in following up, and my time has expired.

The CHAIRMAN. I thank my colleague, and I did not mean to interrupt. We are trying to deal with our clock.

Let's see. Senator Cassidy is next, and then we will have Senator Warner and Senator Warren.

Senator CASSIDY. If I ask questions and say things that you have already answered, I have just been running and dodging protestors. So the question is, how do we regulate this? And I have had three options presented to me, and I would like your kind of take on it—or maybe four.

One, have AI validate AI. And someone says, you've got to be kidding me. AI cannot validate AI because it depends upon making sure that they have the same data sets upon which they are trained and are equally robust.

Second, let's do what we did with the Sherman Antitrust Act. Just give the Court some sort of barriers, guidelines. Okay, courts, you have to stay within this. We cannot imagine what is going to happen in 100 years; courts, figure it out. But then someone said, you've got to be kidding me. You have to understand something before you can have a court ruling, and you think some judge is going to understand this?

It was a lawyer who told me that. I am a gastroenterologist, so I am just going to take his kind of minimization of the judge's ability to comprehend.

The third is to say, if we are talking about health care, to go, for example, to the American College of Cardiology and say we want you to be the third party that will validate that everything out here is going to be okay for a cardiology patient.

Now, since a cardiology patient with congestive heart failure can have kidney disease and can have diabetes and hypertension and be at risk of stroke—and by the way, be HIV positive—it almost seems like you end up with the whole shebang, even though you think you are going to limit it.

But actually, that is the one that seems most valid to me, because you actually have subject matter expertise kind of penetrating there. Now, I have given you that context. Clearly, I am meditating on this before I go to bed. Let's start with you, Dr.

Baicker, and then just work down, and please be tight with your answers.

Dr. BAICKER. Thank you. I think it is crucially important that humans be involved in validating the outcomes, but I hope that there is an opportunity for public-private partnership, and also for consortia of private entities, to have multiple views on the same data set in different contexts and have consensus across them, in which algorithms are actually achieving the goals that they say that they are achieving.

Senator CASSIDY. Practically, that seems like that is going to be incredibly cumbersome, and this seems like a really agile field that is going to race ahead.

Dr. BAICKER. But that is why I think having a fixed structure is less likely to be agile than having multiple parties able to weigh in on whether the outcome is successful.

Senator CASSIDY. Now some of this is proprietary, so would that kind of defeat it if it is proprietary?

Dr. BAICKER. I think there needs to be a regulation that ensures access to the outcomes and to the crucial inputs of the data, so that third parties can see what is happening.

Senator CASSIDY. Okay.

Dr. Obermeyer?

Dr. OBERMEYER. I think there is a nice analogy to how the FDA regulates drugs. When we regulate drugs, the FDA requires that the manufacturer specify the primary outcome—what is the health thing that this drug is affecting—and it requires them to show that the drug affects that outcome in a rigorous way, usually via randomized trial.

With algorithms, I think we should know what the algorithm is predicting, and then I think we should rigorously evaluate how that algorithm is doing its job in a new data set that it has never seen before, which is a key part of how machine learning gets evaluated everywhere. I think that that—

Senator CASSIDY. Now that would presuppose, though, that every time you update it, you have to do retesting.

Dr. OBERMEYER. So I think that there are ways to put in structures for continuous monitoring so it only requires the algorithm to produce its—

Senator CASSIDY. I am running out of time.

Dr. Mello?

Dr. MELLO. I think we have learned that there are models that can usefully evaluate some aspects of other models. For example, the developers at Stanford released a model called APLUS that evaluates algorithmic bias as an open source tool. That is great.

But there are so many other aspects of safe and responsible use of AI that do require human review. I think physicians can be very helpful in spotting issues that will happen when physicians interact with those tools. But generally, I do not think medical professional societies are fully up to the task of vetting some of the issues that arise.

Senator CASSIDY. Then who would be?

Dr. MELLO. I think you are talking about creating what we might call assurance labs, that bring together data scientists with emphasis—

Senator CASSIDY. Okay.

Dr. Sendak?

Dr. SENDAK. In silico testing in a third-party validation environment is necessary but not sufficient. And you know as a practicing physician, different environments are different. Every health-care organization needs to be able to locally govern AI, and that means——

Senator CASSIDY. That seems impossible.

Dr. SENDAK [continuing]. They locally govern quality of care and health-care services——

Senator CASSIDY. I am going to tell you, the small little rural hospital isn't going to do it.

Mr. Shen?

Mr. SHEN. Yes. I think, Senator, it is very important here to really distinguish between FDA-regulated AI and other types of AI like generative AI. I think FDA-regulated AI, again, goes through a regulatory approval process with the FDA where we have to declare the intent of these AI algorithms and what is the clinical use——

Senator CASSIDY. Like clinical support. For example, how do you deal with a heart failure patient?

Mr. SHEN. Exactly. So that framework actually is sufficient to support AI innovation, and we can support kind of this continued flexibility in that approval process to look at how clinical AI can be adopted going forward.

Senator CASSIDY. Thank you. I yield.

The CHAIRMAN. I thank my colleague. We can get both of our remaining Senators in before the noon deadline.

Senator WARNER. Thank you, Mr. Chairman. Thank you all for joining us.

I just want to follow up with Senator Cassidy. Bill? I just want to say one of the things I think we have to grapple with as well is intent, you know? I sit on the Banking Committee, where we are seeing AI tools that might have this notion of “go make money,” but then use illicit or illegal ways to do that.

I think there is a question around intent, maybe to have an AI tool that deals with heart disease, but if you do not have all these other inputs, it could be a challenge.

Dr. Sendak, I want to start with you. I am very interested in your work on developing model disclosures for AI. But I worry—I was a tech guy 20 years before I got into this business—that there is such a first-mover advantage, and if we do not get that disclosure form right—and I know you have talked about using even maybe the nutrition guide tools.

But so, on this kind of model—I think it is an interesting model. You have it going, at least in your testimony I think, to clinical end users. But I do not think you weigh in on whether actually consumers should be getting the same kind of disclosure. And then, how do we make sure that it is actually intelligent, if you do think it—or is there a concern with giving it to the actual patients themselves?

Dr. SENDAK. It is great to see that people actually read papers, because when you write them, sometimes it is like, is anyone going to read this besides your mom?

Senator WARNER. In full disclosure, I have staff that have read everything you have all said. I got the photo op, though. [Laughter.]

Dr. SENDAK. So I want to touch on a few things. You talked about first-mover advantage and the concern with disclosing what is proprietary information. I want to emphasize that—I put this in the testimony—we have incubated four companies.

There are companies that have licensed the algorithms that we have built that they are now commercializing. They are validating them in other settings, and these algorithms all have model facts labels in our publications. So there is a balance between disclosing information about an algorithm and retaining what is needed for commercial rights, to be able to commercialize products.

And then the second point is, who is this for? And I want to emphasize there too, we built that document that you showed for a specific use case, which is for my clinicians, so that they know how these systems work.

That is not the set of information that an organizational governance body needs to be examining to make a procurement decision. That is not the information that a governance committee needs to look at maybe a few times a year, to reconfirm use or look at monitoring. So there is a set of artifacts that need to be developed at different points in time.

Senator WARNER. But again, you talked about the clinical end users. You did not talk about the notion of, should there be some form of this that goes to the patient, and how do we use Medicare as a forcing mechanism to give some level of standardization to a level that even a clinical end user may not understand, let alone an actual patient.

Dr. SENDAK. The only example I am aware of of a public AI inventory by a health-care provider is the VA. We are working through Health AI Partnership to build what we are calling the Community Transparency Registry, where we are going to do one use case—and we will use Duke as an example—where we interview community members, do focus groups, solicit concerns and questions, and develop documentation for the public.

I will say, that is a huge gap. It does not exist today, and it should be part of—

Senator WARNER. I want to get my last question in, and Senator Warren's been very patient. So, Dr. Mello, I just—in your testimony, you talked about automation bias, how we are not going to default to thinking the AI model is always right.

Again, I am interested in how we use Medicare as a tool, and potentially even, I think we do have to raise—I mean, as Senator Cassidy said, you know—some level of liability issues. And I know lots of folks have already talked about bias and hallucination.

I just want to talk to that a little bit more, in terms of how we deal with this automation bias.

Dr. MELLO. Well, I think creating standards—for example, Medicare conditions of participation that require organizations to have given hard evaluation to how the humans interact with the model under the constraints and incentives that they have within that organization—is going to be the key to that.

Senator WARNER. Last, and I just asked this, and I would love to get feedback so I can get—you do not have to answer it. I am chair of the Intel Committee, and Senator Wyden is a current member of that committee.

We spend a lot of time on cyber. There is a whole new way around AI tools, with prompt inquiry and others, that do not really fit neatly into the bucket of cyber, but there are ways that we can manipulate these models in a huge way. I would welcome your input on how we do that.

Thank you, Mr. Chairman.

The CHAIRMAN. I thank my colleague, and I am also interested in that last question of Senator Warner's.

Okay; Senator Warren?

Senator WARREN. All right; thank you.

So, over 31 million Americans are enrolled in Medicare Advantage, or MA, the program that allows private for-profit insurers to offer Medicare coverage. Now, under Federal law, these private insurers are required to cover all Medicare Part A and Part B services.

But in recent years, government watchdogs have found that private insurers are routinely delaying and denying care because doing so boosts their profits. In 2019, the Health and Human Services Inspector General found that nearly one in five payment denials by insurers in MA violated Medicare coverage rules, meaning seniors were unlawfully denied access to services that they were, by law, entitled to.

Some of the largest insurers that offer MA are now relying on flawed artificial intelligence tools like predictive algorithms to scale up their efforts to deny coverage to seniors. These algorithms sift through millions of medical records to determine the level of patient need that the algorithm thinks they need.

So, Professor Mello, you are an expert on AI and health policy. Let us start with an easy question. Does Federal law require all insurance companies to follow Medicare coverage guidelines, even if they are using AI algorithms to determine coverage?

Dr. MELLO. Yes.

Senator WARREN. Yes. So the law is not suspended just because you used AI. I just want to underscore that, because it is clear that these companies are not playing by the rules. So, take UnitedHealthcare, which covers more beneficiaries in MA than any other insurance company.

In 2020, UnitedHealthcare bought NaviHealth, a company that sells its AI services to insurance companies in order to help them make these coverage decisions. Last year, an investigation revealed that UnitedHealthcare had pressured employees to strictly follow NaviHealth's algorithm determinations, leading these human beings to systematically deny care at skilled nursing facilities, even when those decisions were against doctor's orders.

Dr. Obermeyer, you have conducted extensive research on how algorithms are used in health care. Can you talk just a little bit about the dangers that come from solely relying on AI algorithms to make these coverage decisions?

Dr. OBERMEYER. In the case you mentioned, like in all other cases, AI learns from historical data. So, it trawls through those

millions of records and it sees, for example, that there are some privileged people with great insurance who probably stay in nursing homes for longer than they should, and there are also vulnerable underinsured people who are often kicked out too early.

Rather than undoing that problem, the AI reinforces it and encodes it as policy, and I think that is very contrary to the spirit of medical utilization review. It is also a huge missed opportunity, because I think well-designed AI could do much better.

It could look at the patient's X-ray; it could look at the public transportation in their neighborhood; it could look at the layout of their house, and integrate all those things into a far better judgment than a doctor is able to make about who needs to be in that nursing home and who does not.

Senator WARREN. So it is a really important point that you make about how it takes the bad information and accelerates—or the information that tells us about bad practices. You know, according to the investigation from the Inspector General, for some seniors, these AI denials led to “amputations, fast-spreading cancers, and other devastating diagnoses.”

I appreciate that CMS has now finalized a rule to increase transparency requirements on insurers' AI systems and require doctors to verify coverage decisions. But I think we need a lot more to protect patients here. So, if I could come back to you, Professor Mello.

In addition to the rule that the agency has just finalized, what measures do you think that CMS should take to ensure that private insurers are not leveraging AI tools to unlawfully deny care?

Dr. MELLO. Thank you for the question. I think it is really important to look to see whether they are given the incentives involved, and I was very heartened to see in the FAQs released this week on that final rule, CMS plans to beef up its audits in 2024 and specifically look at these denials. That seems extremely important.

But beyond that, I think additional clarification is needed to the plans about what it means to use algorithms properly or improperly. For example, for electronic health records, it did not just say “make meaningful use of those records”; it laid out standards for what meaningful use was.

Senator WARREN. So I think the point here is, we need guard rails, and without significant guard rails in place, these algorithms, as you put it, Dr. Obermeyer, are going to accelerate the problems that we've got and pad private insurers' profits, which gives them even more incentive to use AI in this way.

Until CMS can verify that AI algorithms reliably adhere to Medicare coverage standards by law, then my view on this is, CMS should prohibit insurance companies from using them in their MA plans for coverage decisions. They have to prove they work before they put them in place. Thank you.

Thank you, Mr. Chairman.

The CHAIRMAN. I thank my colleague, and I think Senator Warren's points are very important. And the reality is, not only do we have jurisdiction here, but what is striking—and I have always thought this since my Gray Panther days—is that Medicare is the flagship program in America, and when Medicare does something, everybody in the private sector picks up on it.

So these are very important questions, and I appreciate Senator Warren asking them. She focused on Medicare. I am going to use the word that has not come up over the last 2½ hours, and that is Medicaid. And so here is where we are.

We know Medicaid is a lifeline for millions of the most vulnerable people in the country, and how Medicaid determines eligibility, or the amount of a benefit, is of critical importance. Millions of families and children count on the program for affordable health care.

Increasingly, the States are using algorithms to determine eligibility and benefit-level determinations instead of trained employees, certainly a prospect for distressing outcomes. For example, Tammy Dobbs, an Arkansas woman living with cerebral palsy, had her home care hours suddenly cut by 40 percent after an algorithm was used to redetermine her benefit allotment, leaving her in a state of confusion and panic.

So we have a lot to do to make sure that we shore up the rules to protect people in these health-care plans, because they are particularly vulnerable. And in addition to the States, I would like to address a question to you, Dr. Mello, and you, Dr. Obermeyer, about insurance company discrimination in Medicaid managed care plans.

We are staying on the subject of Medicaid as we wrap up, because those are the folks who are most vulnerable. And something strikes me as fundamentally unfair—and we have been at it for close to 2½ hours—and it never has come up.

So the committee is doing our own investigation into these Medicaid managed care plans, and our investigation includes these issues with respect to insurance company processes that can harm people on Medicaid.

So, Dr. Mello and Dr. Obermeyer, you have some of the best investigators on the planet sitting right behind me here at the dais. What would you tell them they ought to be looking at, in order to protect Medicaid patients, through our inquiry, our investigation that will incorporate these AI issues? What would you tell them to be looking for?

Dr. MELLO. Well, first, thank you so much for undertaking that critically important work. I am even more concerned about those beneficiaries than I am about Medicare beneficiaries, because we know, for a variety of reasons—

The CHAIRMAN. Three cheers for you.

Dr. MELLO. Thank you. They have difficulty challenging those decisions. So, the number one piece of advice I would give is, do not look at grievance appeal and overturn rates as a good indicator of humans in the loop correcting errors, because even if we see—as we have with some other health plans in Medicare Advantage—90 percent overturn rates of these initially mistaken algorithmic decisions, this group of enrollees does not appeal in force. Again, they just simply do not have the social capital to do that.

So, second, I think we want to get a handle on which algorithms are being used and for what purpose; and then third, really look under the hood at the kinds of incentives and constraints the front-line reviewers have, the amount of power they actually have to

overturn those decisions, by talking with them, confidentially, if possible, about, again, what will happen to them if they say “no.”

The CHAIRMAN. The staff was writing very quickly, so you connected.

Dr. Obermeyer?

Dr. OBERMEYER. I want to draw attention to one key difference between Medicare and Medicaid that I think really hurts our ability to understand both the promises and the pitfalls of AI for Medicaid, which is that unlike Medicare, there is no one place you can go to look for data. You have to pool all of the data together from individual States, and for researchers, that is actually really, really difficult, and it prevents people from doing the kind of work that I think is really needed, that I have done around auditing these algorithms.

All that said, I think Medicaid—like everywhere else, we have the same problems of a lot of people who are currently ineligible but who *are* ineligible, and in the people who are covered, there is fraud and abuse. I think AI has a lot of potential to fix both of those problems.

I have seen social workers in the ER struggle to understand who is eligible and who is not. I think there is a huge opportunity for using AI to better structure the eligibility decisions. I also think there is a huge opportunity to use AI to identify exactly those abuses that are happening in these Medicaid managed care plans.

All of that requires access to the data and the ability to do that work.

The CHAIRMAN. We have been at it a long time. You all have been wonderfully patient. This has been a superb hearing. Not a bad point or frankly, from my colleagues, not a bad question in the house, and we have a lot to do.

I will just close by way of saying that one of the other aspects of this that the Senate staff is going to have to get its arms around is, this is very different than what we were tackling in the 1990s, when we were talking about platforms and things like that.

You know today, we are generating a lot of content, and we need really strong guard rails. Those really strong guard rails are just fundamental, and as we kind of march into this kind of new era, we are going to have to have good counsel from all of you. And also, you are going to get some questions for the record. They are due by 5 p.m. on February 15th.

But we want to thank you all, and I look forward to seeing you some time soon, Dr. Mello, on the Stanford campus. I have lots of friends, and you may have run into one of my classmates, Dr. Linda Shortliffe, who taught at the Stanford Medical School—and a lot of other friends.

So, I appreciate all of you, and with that, we are adjourned.

[Whereupon, at 12:11 p.m., the hearing was concluded.]

APPENDIX

ADDITIONAL MATERIAL SUBMITTED FOR THE RECORD

PREPARED STATEMENT OF KATHERINE BAICKER, PH.D.,
PROVOST, UNIVERSITY OF CHICAGO

HARNESSING THE OPPORTUNITIES OF ARTIFICIAL INTELLIGENCE TO TRANSFORM THE HEALTH-CARE SYSTEM

My name is Katherine Baicker, and I am provost of the University of Chicago and a health economics researcher. I would like to thank Senator Wyden, Senator Crapo, and the distinguished members of the committee for giving me the opportunity to speak today about artificial intelligence and health care. I serve on a number of boards and advisory panels but am presenting only my own views. This statement draws on several pieces I have written in this area, as well as research conducted by many others, including some on this panel.

AI holds enormous promise to improve not only the delivery of care, but the effectiveness and sustainability of the health-care system.¹ In addition to improving the quality of care for individual patients, AI can also help us refine the way we pay for health care, focusing resources on patients who would benefit the most. Traditional economics would suggest that payments for health care are the key driver of utilization and value. But it is clear that more nuanced tools are needed. AI tools can improve both clinical decision-making and health care financing.

Of course, AI's promise will only be realized with sufficient large-scale investments and with safeguards to mitigate potential risks. For patient care, AI provides a set of tools to better determine the best course of treatment—but, like other tools, they require rigorous testing in relevant context to judge their utility. For the health-care system, AI provides a potentially transformational avenue to support sustainable investment in long-term population health—but requires public policy guard rails to ensure that all patients can benefit.

AI Tools to Improve Decision-Making and Quality

Improving the quality of care that patients receive is not just about increasing access to new and existing therapies. There is ample evidence that our health-care system has both overuse and underuse of care, leading to worse patient outcomes and unnecessary financial strain. AI tools can help clinicians do better—not just do more or do less. A recent study, for example, deployed a machine learning algorithm to examine how doctors test for acute coronary syndromes (ACS) in the emergency department (ED), and found both overtesting of patients with very low risk (who were extremely unlikely to benefit from the test) and undertesting of patients at high risk (for whom the test would have had high potential to avert severe harms to health).² Reallocating low-value tests to high-risk untested patients could save lives and be extremely cost-effective across a range of care.³ A study of CT pulmonary angiography in national ED visits found similarly large-scale overuse and

¹I focus here primarily on the use of predictive AI, though of course there are many opportunities and challenges with the emerging use of generative AI.

²Mullainathan S, Obermeyer Z. Diagnosing Physician Error: A Machine Learning Approach to Low-Value Health Care. *Quarterly Journal of Economics*. 2022;137(2):1–51.

³Baicker, K and Obermeyer, Z. Overuse and Underuse of Health Care: New Insights from Economics and Machine Learning, *JAMA Health Forum*, 3(2), February 17, 2022.

underuse.⁴ Another study showed that variation in radiologists' diagnostic skill drove over- and underdiagnosis of pneumonia.⁵

AI offers the opportunity to draw in much more information than physicians alone can. In the study of ACS above, researchers found that physicians often focus on a small number of salient variables. Machine learning algorithms can use a much broader set of information to capture the richness of individual patients' histories and conditions. Much of the opportunity lies in bringing together multiple types of data and being able to merge data across silos. Public policy is important here, both because there are limited private incentives to collect and share data that might improve care and because there are real and serious risks to patient privacy and data security. The right information infrastructure with safeguards can not only enable the generation of algorithms to assist clinicians in delivering the right care to their patients but can accelerate discovery of new treatments and modes of care.

But algorithms must supplement, not replace, physicians in care decisions. Physicians can draw in information that algorithms alone can't, including the results of real-time patient exams.⁶ It is crucial that algorithms be tested and validated in relevant contexts, just like any other medical intervention.⁷ The way that information is presented to physicians and integrated into the clinical flow is crucial to improving patient care. And the value that the tools generate should be assessed in terms of real-world improvements in patient outcomes and the efficiency of the resources used. Similarly, algorithms can perpetuate biases present in the care patterns captured in the data used to train the algorithm.⁸ Interrogating the context from which the new information is drawn and the setting in which it will be deployed is crucial to ensuring that incorporating the information benefits all patients.

Improving Delivery and Access at Scale

In addition to improving the quality of care, AI can also help us refine the way we pay for health care, focusing resources on patients who would benefit the most. Economics suggests that payments for health care are the key driver of overuse and underuse—that we see overuse when we pay too much, and underuse when we pay too little. But the fact that we often see both overuse and underuse in the same payment system indicates that more nuanced tools are needed. Aligning payments with the health value that they produce for patients can foster higher-value use of care and minimize spending on care of questionable benefit.⁹ AI can enhance our ability to design and implement value-based insurance and innovative payment systems.^{10, 11}

Of course, coverage of high-value treatments does not necessarily reduce spending. Some treatments, such as childhood immunization and counseling adults about low-dose aspirin to prevent coronary heart disease, may improve health and save money; but most high-value treatments that are highly cost-effective still increase spending.¹² AI can improve our ability to target treatments to the patients who are most likely to benefit from them, allowing coverage of more treatments while promoting affordability of premiums and sustainability of public programs.¹³ Using predictive AI to help identify the patients with the greatest likely benefits can also be a tool to better target insurance coverage expansions.¹⁴

⁴Abaluck J, Agha L, Kabrhel C, Raja A, Venkatesh A. The Determinants of Productivity in Medical Testing: Intensity and Allocation of Care. *American Economic Review*. 2016;106(12):3730–3764.

⁵Chan DC, Gentzkow M, Yu C. Selection with Variation in Diagnostic Skill: Evidence from Radiologists. *The Quarterly Journal of Economics*. Vol. 137 no. 2, 2022.

⁶Agarwal, Moehring, Rajpurkar, and Salz. Combining Human Expertise with Artificial Intelligence: Experimental Evidence from Radiology. NBER Working Paper No. 31422, July 2023.

⁷Shah, Halamka, et al. A Nationwide Network of Health AI Assurance Laboratories. *JAMA*. 2024;331(3):245–249.

⁸Obermeyer, Powers, Vogeli, and Mullainathan. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*. Vol. 366, Issue 6464, pp. 447–453, October 2019.

⁹Baicker, Mullainathan, and Schwartzstein. Behavioral Hazard in Health Insurance, *The Quarterly Journal of Economics* (2015), 1623–1667.

¹⁰Chernew, Rosen, and Fendrick. Value-Based Insurance Design. *Health Affairs*, 26 (2007), w195–w203.

¹¹Baicker, Chernew. Alternative Alternative Payment Models, *JAMA Internal Medicine*, Vol. 177, no. 2, pp 222–223, February 1, 2017.

¹²Neumann PJ, Cohen JT. Cost savings and cost-effectiveness of clinical preventive care. *Synth Proj Res Synth Rep*. 2009;(18):48508.

¹³Baicker, Chandra. Uncomfortable Arithmetic—Whom to Cover Versus What to Cover, *New England Journal of Medicine*, 10.1056/nejmp0911074, Vol. 362, no. 2, 95–97, January 14, 2010.

¹⁴Goto, Inoue, Osawa, Baicker, Fleming, Tsugawa. Machine Learning Detects Heterogeneous Effects of Medicaid Coverage on Depression, *American Journal of Epidemiology*, forthcoming.

Improving individual patient care and affordability is in itself of enormous benefit, but AI also opens up new opportunities at broader scale, unlocking potential innovation in population health management.¹⁵ In our system, payment for care through insurance coverage is a key driver of health care innovation. Insurance plans have latitude about what care they cover, subject to regulation. Enrollees exert some pressure to cover care that they value, but they may not be able to discern the generosity or quality of coverage until they are sick and need specialized care, and the health benefits of preventive care or disease management may not be evident for many years—which may make plans with less coverage and commensurately lower premiums more appealing, affecting health as well as costs.¹⁶ Similarly, employers choosing which plans to offer employees are less likely to internalize the benefits of long-term health because of employee turnover, reducing the incentives to offer plans that invest in care that may only generate improved health (and potentially lower costs) many years later.

AI offers a new way to counteract this disconnect by using better predictions to incentivize coverage of care that improves patient outcomes in the long run. Machine learning applied to data spanning imaging, bloodwork, longitudinal health care use, diagnoses, and more can provide much better information about how future health risks and health-care needs are likely to evolve at the individual and population levels. This information can be used to shape payments to insurers, paying more to those who improve future health prospects for their enrollees and less to those who don't. "Risk adjustment" is already a vital mechanism for mitigating insurers' incentives to enroll only the healthiest patients. Similarly, AI-informed risk adjustment applied to the outcomes that matter most for patients could mitigate insurers' incentives to limit coverage of treatments with short-term costs and long-term health benefits, as well as provide additional information about the quality of care.¹⁷ This would improve access to beneficial care and also foster medical innovation. The absence of better real-time measures of health risk improvement has hindered the development of novel disease management, for example. The development of new health markers through machine learning could unlock new markets for disease management—and thus new modes of addressing serious chronic health conditions.

Predictive AI thus offers much promise at the individual and system level to drive medical innovation, increase the quality of care, and improve patient outcomes—and do so in a way that maintains access and affordability through a focus on high-value use. But that potential hinges on investment in shared data infrastructure and, crucially, on the development of systems of patient protections and algorithm testing and validation that engender trust and ensure broad and equitable patient benefits.

I thank you again for this opportunity and look forward to answering any questions you may have.

QUESTIONS SUBMITTED FOR THE RECORD TO KATHERINE BAICKER, PH.D.

QUESTION SUBMITTED BY HON. MARIA CANTWELL

Question. One of the few sectors where we've successfully implemented meaningful privacy protections is health care, with the enactment of the Health Insurance Portability and Accountability Act (HIPAA). As we continue struggling to bring effective privacy protections to consumers in all sectors of the economy, I'm concerned that with AI in health care we may move backwards in terms of privacy. As we know, AI runs on data—lots of data. And the most valuable data that exists is health-care data. The financial incentive to train AI models on sensitive health data is enormous, but we must also be aware of accidental abuses. For example, we know that Large Language Models tend to "leak" data—including the data they are trained on, as well as data they process during normal use. These mistakes, whether intentional or not, could have enormous consequences for certain people, such as people who travel to another State to get abortions, or those who use menstrual cycle tracker apps that collect sensitive health data. The leaked data could potentially be used to prosecute these patients who are just trying to get reproductive

¹⁵ Baicker, Chandra. Investing in Long-Term Health, *JAMA Health Forum*, February 2024.

¹⁶ Brot-Goldberg ZC, Chandra A, Handel BR, Kolstad JT. What does a deductible do? The impact of cost-sharing on health care prices, quantities, and spending dynamics. *The Quarterly Journal of Economics*. 2017; 132 (3): 1261–1318.

¹⁷ Obermeyer, Powers, Vogeli, and Mullainathan. *Op cit*.

care. That's a critical issue in Washington State, where the number of out-of-state patients seeking abortion care increased by 46 percent in 2022. We should not be compromising the safety of technology users just so AI can collect data and improve algorithms.

What threats does AI pose to privacy in the health-care sector and how often does data get accidentally leaked? How can we improve parameters on data security so that laws like HIPAA can be updated to meet the security requirements in the age of artificial intelligence?

Answer. Fostering responsible use of broad databases has the potential to meaningfully improve patient outcomes, as AI tools can draw on much more information than physicians alone can. Machine learning algorithms can use a much broader set of information to capture the richness of individual patients' histories and conditions and supplement (but not replace) physician decision-making. Much of the opportunity lies in bringing together multiple types of data and being able to merge data across silos, but there are limited private incentives to collect and share data that might improve care. The right information infrastructure can not only enable the generation of algorithms to assist clinicians in delivering the right care to their patients, but can accelerate discovery of new treatments and modes of care.

To achieve these benefits, it is crucial that algorithms be tested and validated in relevant contexts. The way that information is presented to physicians and integrated into the clinical flow is a crucial determinant of effectiveness, and the value that the tools generate should be assessed in terms of real-world improvements in patient outcomes. The development of systems of patient protections and algorithm testing and validation standards that engender trust and ensure broad and equitable patient benefits is vital.

QUESTIONS SUBMITTED BY HON. SHELDON WHITEHOUSE

Question. Rhode Island has several health-care enterprises that we are very proud of, including the State-wide Rhode Island All-Payer Claims Database (RI APCD).

What role can AI play with regard to maximizing the utility of a robust, All-Payer Claims Database?

Answer. In addition to improving the quality of care, AI can also help us refine the way we pay for health care, focusing resources on patients who would benefit the most. The fact that we often see both overuse and underuse in the same payment system indicates that more nuanced tools are needed. Aligning payments with the health value that they produce for patients can foster use of care with high patient health benefits while minimizing spending on care that does little to improve health. AI can enhance our ability to design and implement value-based insurance and innovative payment systems, promoting affordability of premiums and sustainability of public programs. Transparency for both patients and providers can amplify the effectiveness of value-based payments and insurance design. Using predictive AI to help identify the patients with the greatest likely benefits can also be a tool to better target insurance coverage expansions. Robust databases are key enablers of unlocking this value.

Question. Are you working in the development of AI with some of the medical specialty organizations (orthopedics, cardiologists, etc.)? Are specialty organizations a useful place for benchmarking, approval, or accrediting best AI practices? And if they are, do you have any good examples of a medical specialty association that is being particularly forward and helpful at looking for the best uses of AI within the specialty?

Rhode Island has two, unusually good Accountable Care Organizations (ACOs). In what ways could artificial intelligence be used to support the ACO program?

Answer. Predictive AI offers much promise at the individual and system level to drive medical innovation, increase the quality of care, and improve patient outcomes. There is, however, a demonstrated risk that algorithms can perpetuate biases present in the care patterns captured in the data used to train the algorithm. Interrogating the context from which the new information is drawn and the setting in which it will be deployed is crucial to ensuring that incorporating the information benefits all patients.

Improving the quality of care that patients receive is not just about increasing access to existing therapies. Payment for care through insurance coverage is a key

driver of health-care innovation and the creation of new therapies. The absence of better real-time measures of health risk improvement has hindered the development of novel disease management, for example. AI-informed risk adjustment applied to the outcomes that matter most for patients could mitigate insurers' incentives to limit coverage of treatments with short-term costs and long-term health benefits, as well as provide additional information about the quality of care. This would improve access to beneficial care and also foster medical innovation. The development of new health markers through machine learning could unlock new markets for disease management—and thus new modes of addressing serious chronic health conditions through ACOs or other structures, particularly for underserved populations. I myself am not currently working with any medical specialty organizations.

PREPARED STATEMENT OF HON. MIKE CRAPO,
A U.S. SENATOR FROM IDAHO

Artificial intelligence, or AI, offers seemingly endless applications and opportunities for every sector of American society, including health care. From disease detection and diagnosis to advanced imaging, hundreds of AI-enabled devices have come to market in recent years, providing critical tools for front-line clinicians. AI can also streamline and simplify taxing administrative tasks, which continue to pull providers away from patient care, imposing high costs and eroding outcomes.

Every week, new breakthroughs in AI technology come to light, presenting countless use-cases to drive health-care improvements, from more efficient drug discovery to more informed clinical decisions.

At the same time, as with any emerging technology, AI also raises new questions and risks. Responsible and ethical deployment—with appropriate safeguards for privacy and security—will prove crucial to building patient trust and ensuring effective results.

As we weigh the promise of AI in the health-care context, our committee will play a critical role, given our jurisdiction over a range of Federal health programs.

Today's hearing will help us to strike a patient-centered and fiscally responsible approach to AI-enabled tools within these and other Federal programs. Rather than legislate first and ask questions later, Congress needs to build a better understanding of how AI operates in specific contexts, including health care. Our witnesses will identify areas where current policies can adapt effectively, as well as where they fall short. These discussions can lend valuable insights into the real-world promise of innovation for patients, along with the importance of reliability, transparency, and adaptability—especially as policymakers navigate a rapidly changing landscape.

As game-changing AI-enabled devices and other technologies emerge, Medicare coverage and payment policies must keep pace. Otherwise, access gaps for seniors will widen, care quality will suffer, and the innovation pipeline will shrink.

I look forward to engaging with our witnesses, as well as my colleagues in both chambers, to ensure we address the regulatory hurdles and pervasive uncertainty that too often confront older Americans and those living with disabilities. Our Federal programs should expedite, rather than undermine, access to medical breakthroughs, including those enabled by AI.

In order to build trust in high-quality innovations, we also need appropriately targeted oversight and analysis, which can help to inform future policy initiatives. In the case of algorithms applied to expedite utilization management, for instance, improper denials or delays of needed services warrant government scrutiny. In other contexts, AI-related challenges stem largely from insufficient provider experience or education with cutting-edge tools.

In the latter case, rather than rely on a top-down approach, Federal agencies should leverage expertise, both in-house and external, to scale up outreach and technical assistance, particularly for lower-resource providers and sites, including in rural areas. Collaborative public-private partnerships have the potential to drive responsible deployment across the country while generating vital data and promoting robust privacy protections for patients.

Beneficiary trust and clinician uptake will demand a transparent, but also risk-based, approach to the use of AI-assisted tools, tailored to the vast differences among diverse technologies and use-cases. Along these lines, AI highlights the need

for adaptability. One-size-fits-all, overly rigid, and unduly bureaucratic laws and regulations risk stifling lifesaving advances and becoming outdated before they are even codified. Only with a sector-specific approach, focused squarely on our jurisdiction—and informed by experts and stakeholders from across the field—can we pursue responsible and responsive policies that improve access, enhance care quality, and reduce care costs.

Thank you to our witnesses for being here today. I look forward to your testimony.

PREPARED STATEMENT OF MICHELLE M. MELLO, JD, PH.D.,¹
PROFESSOR OF HEALTH POLICY AND LAW, STANFORD UNIVERSITY

FACILITATING RESPONSIBLE GOVERNANCE OF HEALTH-CARE AI TOOLS

Chairman Wyden, Ranking Member Crapo, and members of the committee, thank you for the opportunity to speak with you today.

I have the extraordinary privilege of being part of a group of ethicists, data scientists, and physicians at Stanford University—long a leading hub of AI innovation—that is directly involved in governing how health-care AI tools are used in patient care. I have studied patient safety, health care quality regulation, and data ethics for more than 2 decades. I apply that expertise in our team’s evaluations of all AI tools proposed for use in Stanford Health-Care facilities, which care for over 1 million patients per year, and our recommendations about whether and how they can be used safely and effectively.

I would like to share the three most important things we’ve learned so far.

First, while hospitals are starting to recognize the need to vet AI tools before use, **most health-care organizations don’t have robust review processes yet.** Some, like Stanford, have plentiful resources to drawn on; others don’t. All need help. Although as a lawyer I know that more law isn’t always the answer, in this case there is much that Congress could do to help.

Second, to be effective, **governance can’t focus only on the algorithm. It must also encompass how the algorithm is integrated into clinical workflow.** By “workflow,” I mean how physicians, nurses, and other staff interact with each other, the AI tool, the patient, and other systems. Currently, conversations about regulating health-care AI mostly focus on the AI tool itself—for example, is its output biased? How often does it make wrong predictions or misclassify things?

These things matter. But it is equally important to consider how medical professionals will interact with the tool. A key area of inquiry is the expectations placed on physicians and nurses to evaluate whether AI output is accurate for a given patient, given the information readily at hand and the time they will realistically have. For example, large-language models like ChatGPT are employed to compose summaries of clinic visits and doctors’ and nurses’ notes, and to draft replies to patients’ emails. Developers trust that doctors and nurses will carefully edit those drafts before they’re submitted—but will they? Research on human-computer interactions shows that humans are prone to automation bias: we tend to overrely on computerized decision support tools and fail to catch errors and intervene where we should.²

Therefore, regulation and governance should address not only the algorithm, but also how the adopting organization will use and monitor it. To take a simple analogy, if we want to avoid motor vehicle accidents, we can’t just set design standards for cars. Road safety features, driver’s licensing requirements, and rules of the road all play important roles in keeping people safe.

Third, because the success of AI tools depends on the adopting organization’s ability to support them through vetting and monitoring, **the Federal Government should establish standards for organizational readiness and responsibility to use health-care AI tools, as well as for the tools themselves.** As countless

¹Professor of health policy, Department of Health Policy, Stanford University School of Medicine; professor of law, Stanford Law School; affiliate faculty, Stanford Institute for Human-Centered Artificial Intelligence; institute faculty, Freeman Spogli Institute for International Studies, Stanford University.

²Mello MM, Guha N. Understanding liability risk from using healthcare artificial intelligence tools. *N Engl J Med*. 2024;390(3):271–278. Full text available from the author (mmello@law.stanford.edu).

historical examples of medical innovations have shown, having good intentions isn't enough to protect against harm. The community needs some guard rails and guidance.

I believe there is a right and a wrong way to do this. It would be a mistake to enshrine in legislation detailed standards for health-care AI tools and how they can be used. In light of how quickly things are moving in the field, we have to have the humility to acknowledge that we don't know what the best standards will be 2 years from now. Regulation needs to be adaptable or else it will risk irrelevance—or worse, chilling innovation without producing any countervailing benefits. **The wisest course now is for the Federal Government to foster a consensus-building process that brings experts together to create national consensus standards and processes for evaluating proposed uses of AI tools.**

It can also begin requiring that entities regulated by Federal agencies adhere to those standards and processes. Through its operation of and certification processes for Medicare, Medicaid, the Veterans Affairs Health System, and other health programs, Congress and Federal agencies can require that participating hospitals and clinics have a process for vetting any AI tool that affects patient care before deployment and a plan for monitoring it afterwards. As an analogue, the Centers for Medicare and Medicaid Services (CMS) uses The Joint Commission, an independent, not-for-profit organization, to inspect health-care facilities for purposes of certifying their compliance with the Medicare Conditions of Participation. The Joint Commission recently developed a voluntary certification standard for the Responsible Use of Health Data which focuses on how patient data will be used to develop algorithms and pursue other projects. A similar certification could be developed for facilities' use of AI tools.

The initiative currently underway to create a network of “AI assurance labs,”³ and consensus-building collaboratives like the 1,400-member Coalition for Health AI, can be pivotal supports for these facilities. Such initiatives can develop consensus standards, provide technical resources, and perform certain evaluations of AI models, like bias assessments, for organizations that don't have the resources to do it themselves. Adequate funding will be crucial to their success.

It's important to recognize that some aspects of AI review will need to be done locally, as individual health-care organizations are best positioned to identify and address problems that could result from how they embed AI tools within clinical workflow.⁴ Here, too, regulatory requirements can ensure that organizations invest in making it happen—just as the Federal regulations known as “the Common Rule” did for ethical review of human subjects research.

We have developed such a review process at Stanford. For each AI tool proposed for deployment in Stanford hospitals, data scientists evaluate the model for bias and clinical utility. Ethicists interview patients, clinical care providers, and AI tool developers to learn what matters to them and what they're worried about.⁴ We find that with just a small investment of effort, we can spot potential risks, mismatched expectations, and questionable assumptions that we and the AI designers hadn't thought about. In some cases, our recommendations may halt deployment; in others, they strengthen planning for deployment. We designed this process to be scalable and exportable to other organizations.

I will close with one final point from other research I have conducted on AI: **don't forget health insurers.** Just as with health-care organizations, real patient harm can result when insurers use algorithms to make coverage decisions. For instance, members of Congress have expressed concern about Medicare Advantage plans' use of an algorithm marketed by NaviHealth in prior-authorization decisions for post-hospital care for older adults. In theory, human reviewers were making the final calls while merely factoring in the algorithm output; in reality, they had little discretion to overrule the algorithm. This is another illustration of why humans' responses to model output—their incentives and constraints—merit oversight.

CMS recently took the important step of addressing these practices through a Final Rule requiring Medicare Advantage plans to make medical necessity determinations “based on the circumstances of the specific individual . . . as opposed to using an algorithm or software that doesn't account for an individual's circum-

³Shah NH, Halamka JD, Saria S, et al. A nationwide network of health AI assurance laboratories. *JAMA*. 2024;331(3):245–249.

⁴Mello MM, Shah NH, Char DS. President Biden's executive order on artificial intelligence—implications for health care organizations. *JAMA*. 2024;331(7):17–18.

stances.” The final rule further specifies that determinations must be reviewed by a medical professional. But ambiguity remains around what it means to merely “use” algorithms, as opposed to allowing them to drive decisions; and what it means to “account for” individual circumstances or to have algorithm results “reviewed by” a human.⁵ How much freedom must human reviewers have to overrule algorithm recommendations? Must algorithms include information on social determinants of health or patients’ social supports to “account for” individual circumstances? Must insurers disclose the prediction algorithm? Additional clarity about regulators’ expectations would be very helpful.

In summary, Congress can support health-care organizations and health insurers navigating the uncharted territory of AI tools by imposing some guard rails while allowing the rules to evolve with the technology. Specifically, Congress should:

1. Require that health-care organizations have robust processes for determining whether planned uses of AI tools meet certain standards, including undergoing ethical review.
2. Fund a network of AI assurance labs to develop consensus-based standards and ensure that lower-resourced health-care organizations have access to necessary expertise and infrastructure to evaluate AI tools.
3. Require developers of AI tools to disclose information that a consensus-based organization such as an assurance lab determines is essential to evaluating the safety and ethics of AI tools.
4. Work with CMS to provide further guidance to Medicare Advantage plans about permissible and impermissible uses of algorithms in coverage decisions.
5. Ensure that relevant Federal agencies, including but not limited to CMS, the Department of Veterans Affairs, and the Food and Drug Administration, have clear grants of authority to adopt standards for all types of health-care AI and require entities within their purview to adhere to them. Clarity and specificity are essential here, because courts are increasingly holding that on matters of vast social and economic significance like AI, Congress must speak clearly when it intends to give authority to agencies.

Thank you, and I welcome your questions.

QUESTIONS SUBMITTED FOR THE RECORD TO MICHELLE M. MELLO, JD, PH.D.

QUESTIONS SUBMITTED BY HON. RON WYDEN

Question. Several recent reports have highlighted that more people are being denied care with the help of Artificial Intelligence (AI) tools. For people who buy their own health insurance—and are not provided coverage through their employer—one in five in-network claims were denied. The Health and Human Services Inspector General found similarly high denial rates in Medicaid managed care plans.

Do you think one of the reasons for these high denial rates is because AI or other algorithms are supercharging the process?

What should Congress do to strike the balance between innovation and consumer protection?

Answer. The facts behind these highly publicized insurance denials are still being sussed out in litigation and other investigations. Algorithms certainly have the ability to propagate errors on a massive scale, and therefore are rightly a focus of these investigations.

It is critical for Congress and CMS to exercise close oversight of how Medicare Advantage plans and other insurers are using algorithms in coverage decisions, including but not limited to algorithms that use artificial intelligence/machine learning (AI/ML). CMS recently took the important step of addressing these practices through a final rule requiring Medicare Advantage plans to make medical necessity determinations “based on the circumstances of the specific individual . . . as opposed to using an algorithm or software that doesn’t account for an individual’s cir-

⁵ Mello MM, Rose S. Denial—artificial intelligence tools and health insurance coverage decisions. *JAMA Health Forum* (forthcoming; available from the author (mmello@law.stanford.edu)).

cumstances.” The final rule further specifies that determinations must be reviewed by a medical professional.

However, more could be done to reduce ambiguity around what it means to merely “use” algorithms, as opposed to allowing them to drive decisions; what it means to “account for” individual circumstances; and what it means to have algorithm results “reviewed by” a human.¹ For example, a lawsuit against Cigna over its use of a non-AI-driven algorithm to make medical necessity determinations, alleges that the reviewing physicians spent, on average, 1.2 seconds per denied claim. That is human review in name only.

To establish a good balance between facilitating innovation and protecting consumers, Congress and CMS should set concrete standards for insurers’ use of algorithms that provide greater specificity regarding what is and isn’t acceptable.

For example, what must insurers disclose about the factors that algorithms take into account in generating output? Are there certain factors that algorithms for different types of decisions must include or may not include? What kinds of plans should insurers have in place for testing whether the algorithm is performing as well for their patients as the developer claimed it would, based on training data? Are there circumstances in which human reviewers must reach out to the care team and patient or family members before ratifying an algorithmic decision to deny a claim—if so, what are they? CMS should also vigorously pursue its planned audits of health plans’ use of algorithms, as described in its February 2024 guidance document.

Question. My recent oversight efforts into medical privacy have demonstrated that Medicare and Medicaid beneficiaries’ pharmacy records are vulnerable to seizure by law enforcement at the pharmacy counter. Recent news reports suggest that transgender children and adults’ medical records are similarly vulnerable to seizure by health oversight authorities at hospitals and clinics.

What more should Congress do to ensure that patients’ medical data is being protected by health-care entities?

What can, and should, hospitals, clinics, and other providers already be doing on their own to safeguard the privacy of Medicare and Medicaid patients?

Answer. AI models are usually developed on data that don’t include information that directly identifies patients (*i.e.*, “deidentified” data). When a hospital shares fully deidentified data with an external algorithm developer for purposes of creating or improving an algorithm, it doesn’t implicate HIPAA or other information privacy laws. Some patients undoubtedly would object to having their data used by third-party developers in this way, and in the future it will increasingly be possible to identify patients by triangulating putatively deidentified data from different sources. But presently this risk is not high—if data are leaked or hacked, for example, it’s not immediately apparent whose data the unauthorized party is looking at. Imposing restrictions on sharing deidentified data would certainly dampen innovation and keep it concentrated at large companies and medical centers who already possess big datasets of patient information.

Requests by law enforcement for information about specific patients are an entirely different problem. Here, patients are identified, and the party that wants the data usually wants it in order to take some adverse action against them (such as prosecuting patients for obtaining abortion medications or prosecuting doctors for helping patients obtain abortion care). Currently, HIPAA permits but does not require hospitals, pharmacies, and other covered entities to disclose patient data in response to law enforcement requests. They may be inclined to comply because resistance invites further legal tussles.²

Congress and HHS could strengthen patient protections by advancing the HHS proposed rule that would prohibit disclosures of identifiable patient information in proceedings against individuals relating to reproductive health care. The proposed rule limits protection to care that was legal in the State where it was provided; protection could be expanded by eliminating that limitation. Proposals to require a warrant in order to obtain medical records would also strengthen patient protections,

¹ Mello MM, Rose S. Denial—artificial intelligence tools and health insurance coverage decisions. *JAMA Health Forum* (forthcoming on March 7, 2024; to be available at <https://jamanetwork.com/journals/jama-health-forum> or from the authors at mmello@law.stanford.edu).

² Spector-Bagdady K, Mello MM. Protecting the privacy of reproductive health information after the fall of *Roe v Wade*. *JAMA Health Forum*. 2022;3(6):e222656.

though they would not prevent law enforcers from obtaining medical records in every instance. Finally, Congress could extend patient protections by adopting new privacy rules for online health resources such as fertility trackers and other apps, since HIPAA only reaches entities that provide health care and meet other requirements.

QUESTIONS SUBMITTED BY HON. CHUCK GRASSLEY

Question. I'm a strong defender of maintaining and improving access to rural health care.

Can you describe policy considerations we should take into account when promoting and paying for artificial intelligence in rural health care?

Answer. There are two key things to keep in mind when thinking about how to ensure that the benefits of AI and other forms of health IT are equitably shared by rural health facilities and the patients they serve. First, many rural facilities do not have the technical resources (human or computer) to conduct sophisticated evaluations of potential IT tools or implement these tools with extensive monitoring. It is not realistic to expect small, rural health-care facilities to develop a great deal of AI expertise in the near term, even with financial help. There are certainly aspects of responsible AI use that can and should be done locally, such as assessment of how the facility will train nurses and physicians on interacting with an AI tool, and how the tool will be integrated into these practitioners' workflow. But for other aspects of AI evaluation, it makes more sense to have evaluation services provided by larger organizations. The current initiative to create AI "assurance labs" will do so, if Congress ensures it is adequately funded.

Second, many rural facilities do not have extra funds to invest in licensing and implementing AI tools that may improve quality of care but won't generate new revenue or save them money. If Congress wants such tools to be widely adopted, it and HHS should design programs to subsidize adoption. This could be done directly (*e.g.*, through grants or tax credits) or by adjusting reimbursement for particular services in Medicare, Medicaid, and other insurance programs.

We do not want to cultivate a system in which patients in rural areas receive less benefit or are exposed to higher risk from AI tools than patients in better-resourced areas. Nor do we want a system in which patients see surcharges on their medical bills for facilities' use of AI tools, forcing them to pay out of pocket for AI or make tough decisions about whether to opt in to use of a particular AI tool in their care at their expense. Just as with adoption of electronic health records, Federal programs can provide strong financial incentives to assure a different future.

Question. Are you aware of any government oversight entities, such as the Government Accountability Office or Inspector General, using artificial intelligence to stop waste, fraud, and abuse?

If so, how are they using artificial intelligence, and is it effective? If not, how can we deploy artificial intelligence as an oversight tool?

Answer. I do not have the necessary information to provide a response.

QUESTIONS SUBMITTED BY HON. MARIA CANTWELL

Question. One of the few sectors where we've successfully implemented meaningful privacy protections is health care, with the enactment of the Health Insurance Portability and Accountability Act (HIPAA). As we continue struggling to bring effective privacy protections to consumers in all sectors of the economy, I'm concerned that with AI in health care we may move backwards in terms of privacy.

As we know, AI runs on data—lots of data. And the most valuable data that exists is health-care data. The financial incentive to train AI models on sensitive health data is enormous, but we must also be aware of accidental abuses. For example, we know that Large Language Models tend to "leak" data—including the data they are trained on, as well as data they process during normal use.

These mistakes, whether intentional or not, could have enormous consequences for certain people, such as people who travel to another State to get abortions, or those who use menstrual cycle tracker apps that collect sensitive health data. The leaked data could potentially be used to prosecute these patients who are just trying

to get reproductive care. That's a critical issue in Washington State, where the number of out-of-State patients seeking abortion care increased by 46 percent in 2022.

We should not be compromising the safety of technology users just so AI can collect data and improve algorithms.

What threats does AI pose to privacy in the health-care sector, and how often does data get accidentally leaked?

How can we improve parameters on data security so that laws like HIPAA can be updated to meet the security requirements in the age of artificial intelligence?

Answer. AI models are usually developed on data that don't include information that directly identifies patients (*i.e.*, "deidentified" data). When a hospital shares fully deidentified data with an external algorithm developer for purposes of creating or improving an algorithm, it doesn't implicate HIPAA or other information privacy laws. Some patients undoubtedly would object to having their data used by third-party developers in this way, and in the future it will increasingly be possible to identify patients by triangulating putatively deidentified data from different sources. But presently this risk is not high—if data are leaked or hacked, for example, it's not immediately apparent whose data the unauthorized party is looking at. Imposing restrictions on sharing deidentified data would certainly dampen innovation and keep it concentrated at large companies and medical centers who already possess big datasets of patient information.

Requests by law enforcement for information about specific patients are an entirely different problem. Here, patients are identified, and the party that wants the data usually wants it in order to take some adverse action against them (such as prosecuting patients for obtaining abortion medications or prosecuting doctors for helping patients obtain abortion care). Currently, HIPAA permits but does not require hospitals, pharmacies, and other covered entities to disclose patient data in response to law enforcement requests. They may be inclined to comply because resistance invites further legal tussles.³

Congress and HHS could strengthen patient protections by advancing the HHS proposed rule that would prohibit disclosures of identifiable patient information in proceedings against individuals relating to reproductive health care. The proposed rule limits protection to care that was legal in the State where it was provided; protection could be expanded by eliminating that limitation. Proposals to require a warrant in order to obtain medical records would also strengthen patient protections, though they would not prevent law enforcers from obtaining medical records in every instance. Finally, Congress could extend patient protections by adopting new privacy rules for online health resources such as fertility trackers and other apps, since HIPAA only reaches entities that provide health care and meet other requirements.

Question. I think everybody understands that health care is a deeply personal vocation. I don't think we'll ever have machines that can sit up all night holding a dying patient's hand. So we all appreciate how important it is to preserve the human element in health care.

When people talk about AI technology in general, the phrase they use is "a human in the loop"—we need to keep a human in the loop. But in health care, I think we want a lot more than a human "in the loop." We need to be certain that crucial decisions are being made by humans.

It's great that AI can help us make better decisions, but I'm concerned that things might go the other way, where humans merely implement decisions made by machines. We know that humans often tend to believe what a computer tells them. This issue is compounded by the fact that the data that algorithms use can contain harmful biases, which would eventually be amplified to make decisions that could hurt health equity and lead to negative health outcomes for certain population groups. That is why it is so crucial that humans remain the final decision-maker so that we can avoid a situation where the tail is wagging the dog.

How can we be sure, as AI is integrated into our health-care systems, that human judgement will not be subordinated to the judgement of machines?

Answer. We can ensure this by setting standards for, and monitoring, how medical professionals interact with AI tools. Right now, conversations about AI regula-

³Spector-Bagdady K, Mello MM. Protecting the privacy of reproductive health information after the fall of *Roe v Wade*. *JAMA Health Forum*. 2022;3(6):e222656.

tion and governance focus heavily on the algorithms themselves—but equally important is how the adopting organization will use and monitor the algorithms. To take a simple analogy, if we want to avoid motor vehicle accidents, we can't just set design standards for cars. Road safety features, driver's licensing requirements, and rules of the road all play important roles in keeping people safe.

Because the success of AI tools depends on the adopting organization's ability to support them through vetting and monitoring, the Federal Government should establish standards for organizational readiness and responsibility to use health-care AI tools, as well as for the tools themselves. For example, how will hospitals ensure that their clinical staff don't fall victim to automation bias (humans' well-documented tendency to over-rely on computerized decision support tools and fail to catch errors and intervene where we should)? How will they ensure that for each tool implemented, there is a well-conceived plan for monitoring for automation bias and intervening—including stopping use of the tool—if problems are discovered?

Question. Is it possible to weed out harmful biases in AI algorithms? What type of guard rails can we implement to ensure that harmful biases do not form the basis of health-care decisions recommended by artificial intelligence?

Answer. First, it is important to understand that very often, biases in algorithms aren't really about the algorithms—they're about the data used to train them. If particular demographic groups are underrepresented in training datasets, the risk that a model will underperform for those groups is heightened, because the model wasn't given enough information to learn how to get it right. Relatedly, if model developers use suboptimal data as a proxy for the thing they're really trying to predict (for example, using people's health-care utilization to approximate their health-care need), the results will be concomitantly poor. Therefore, a critical component of improving equity by design is improving the richness and representativeness of the datasets available to train them.

To improve data resources, Federal agencies could release more of the data the government currently holds, and allow its use with fewer restrictions and costs. Often, data access is restricted in order to protect marginalized groups, but the restrictions can come at the expense of those very same groups. Congress and HHS could also require more collection and clear labeling of demographic groups in various datasets. For instance, many radiographic images are not marked as belonging to a child versus an adult.

Second, while these steps would be helpful, it is probably not possible to design away all bias. For example, even free-flowing electronic health record data or insurance claims data wouldn't change the fact that some population groups have lower access to health care, so will be underrepresented in those data until our health system improves equity of access.

Because bias is always a risk, it is critical that both algorithm developers and algorithm users test for it. This should be a key part of any Federal standards for algorithms as well as the implementation plans that hospitals and other facilities put together when they want to use algorithms, *i.e.*, what will the hospital do (perhaps in partnership with the developer or another external organization like an assurance lab) to show that, 6 months into using an AI tool, it is performing as well for Black patients as for Caucasian patients?

Finally, the humans using AI tools need to be trained to account for bias when they make decisions. If an algorithm performs better for some groups of patients than others, the physicians, nurses, and facility administrators who use the algorithm need to know that, and to be reminded to think about it when they are deciding how, if at all, to use model output when they make decisions.

QUESTIONS SUBMITTED BY HON. JOHN CORNYN

Question. Last month, HHS's Office of the National Coordinator for Health IT finalized a rule establishing new data standards and reporting requirements for clinicians using ONC-certified health IT. As this technology becomes more prevalent in patient care, we must ensure patient safety but also recognize that smaller or rural providers might not have the same resources to implement these new requirements. For example, while larger health systems may have a designated Chief Technology Officer, smaller practices may have someone who is juggling this and many other roles.

What are the risks of creating new data or IT requirements that smaller health systems do not have the tools to properly implement?

Are there incentives we can use to encourage health-care providers to implement IT best practices?

Answer. There are two key things to keep in mind when thinking about how to ensure that the benefits of AI and other forms of health IT are equitably shared by rural health facilities and the patients they serve. First, many rural facilities do not have the technical resources (human or computer) to conduct sophisticated evaluations of potential IT tools or implement these tools with extensive monitoring. It is not realistic to expect small, rural health-care facilities to develop a great deal of AI expertise in the near term, even with financial help. There are certainly aspects of responsible AI use that can and should be done locally, such as assessment of how the facility will train nurses and physicians on interacting with an AI tool, and how the tool will be integrated into these practitioners' workflow. But for other aspects of AI evaluation, it makes more sense to have evaluation services provided by larger organizations. The current initiative to create AI "assurance labs" will do so, if Congress ensures it is adequately funded.

Second, many rural facilities do not have extra funds to invest in licensing and implementing AI tools that may improve quality of care but won't generate new revenue or save them money. If Congress wants such tools to be widely adopted, it and HHS should design programs to subsidize adoption. This could be done directly (*e.g.*, through grants or tax credits) or by adjusting reimbursement for particular services in Medicare, Medicaid, and other insurance programs.

We do not want to cultivate a system in which patients in rural areas receive less benefit or are exposed to higher risk from AI tools than patients in better-resourced areas. Nor do we want a system in which patients see surcharges on their medical bills for facilities' use of AI tools, forcing them to pay out of pocket for AI or make tough decisions about whether to opt in to use of a particular AI tool in their care at their expense. Just as with adoption of electronic health records, Federal programs can provide strong financial incentives to assure a different future.

QUESTIONS SUBMITTED BY HON. SHERROD BROWN

Question. I am deeply concerned about the many barriers that patients face in accessing timely care, including arduous prior authorization processes.

I introduced bipartisan legislation, the Improving Seniors' Timely Access to Care Act, which would modernize and streamline prior authorization processes in the Medicare Advantage program and improve access to timely care. Many of the policies included in this legislation have been finalized by the Centers for Medicare and Medicaid Services (CMS).

I understand that Artificial Intelligence (AI) could be utilized as a tool to quicken prior authorization processes and increase access to care. While this is a significant step forward in eliminating barriers that delay care to patients, it is important that there are sufficient guardrails in place to prevent people from systematically or inaccurately being denied treatment.

What are some ways that Congress can better oversee the use of AI in the prior authorization process?

How can Congress ensure that AI does not cause a disruption in or denial of care that beneficiaries are entitled to under the law?

Answer. The facts behind these highly publicized insurance denials are still being sussed out in litigation and other investigations. Algorithms certainly have the ability to propagate errors on a massive scale, and therefore are rightly a focus of these investigations.

It is critical for Congress and CMS to exercise close oversight of how Medicare Advantage plans and other insurers are using algorithms in coverage decisions, including but not limited to algorithms that use artificial intelligence/machine learning (AI/ML). CMS recently took the important step of addressing these practices through a final rule requiring Medicare Advantage plans to make medical necessity determinations "based on the circumstances of the specific individual . . . as opposed to using an algorithm or software that doesn't account for an individual's circumstances." The final rule further specifies that determinations must be reviewed by a medical professional.

However, more could be done to reduce ambiguity around what it means to merely “use” algorithms, as opposed to allowing them to drive decisions; what it means to “account for” individual circumstances; and what it means to have algorithm results “reviewed by” a human.⁴ For example, a lawsuit against Cigna over its use of a non-AI-driven algorithm to make medical necessity determinations, alleges that the reviewing physicians spent, on average, 1.2 seconds per denied claim. That is human review in name only.

To establish a good balance between facilitating innovation and protecting consumers, Congress and CMS should set concrete standards for insurers’ use of algorithms that provide greater specificity regarding what is and isn’t acceptable. For example, what must insurers disclose about the factors that algorithms take into account in generating output? Are there certain factors that algorithms for different types of decisions must include or may not include? What kinds of plans should insurers have in place for testing whether the algorithm is performing as well for their patients as the developer claimed it would, based on training data? Are there circumstances in which human reviewers must reach out to the care team and patient or family members before ratifying an algorithmic decision to deny a claim—if so, what are they? CMS should also vigorously pursue its planned audits of health plans’ use of algorithms, as described in its February 2024 guidance document.

QUESTIONS SUBMITTED BY HON. SHELDON WHITEHOUSE

Question. Rhode Island has several health-care enterprises that we are very proud of, including the State-wide Rhode Island All-Payer Claims Database (RI APCD).

What role can AI play with regard to maximizing the utility of a robust, All-Payer Claims Database?

Answer. The potential to develop AI tools has increased demand for patient data, which can be used to train and test AI models. All-payer claims databases are valuable because, relative to other commonly used data sources, they are broader in geographic scope and/or the groups of patients represented. Because they may be more representative of the patient population as a whole, they are especially useful for avoiding the biases that come from training models on narrower groups of patients. Yet, it is important to keep in mind that claims databases represent people’s health-care use, not their health-care needs. Because people face various barriers to accessing care, health-care use isn’t necessarily a good proxy of who needs what health care or could benefit from more health care.

Initiatives to make it easier for researchers and model developers to access all-payer databases—with appropriate data protections—can advance these goals. Policymakers may wish to consider whether access should be equally available to commercial developers and not-for-profit organizations (*e.g.*, academic medical centers that are developing and testing algorithms) or whether for-profit companies’ access to patient data should be more restricted.

Question. Are you working in the development of AI with some of the medical specialty organizations (orthopedics, cardiologists, etc.)? Are specialty organizations a useful place for benchmarking, approval, or accrediting best AI practices? And if they are, do you have any good examples of a medical specialty association that is being particularly forward and helpful at looking for the best uses of AI within the specialty?

Answer. I am not presently working with any physician specialty organizations, but I have recently begun helping the American Hospital Association (AHA) perform an assessment of an AI algorithm designed to improve diagnosis of a serious cardiac condition. Presently, medical professional organizations don’t appear to have the capacity to evaluate AI algorithms themselves. They could, however, be very helpful as consultants to organizations that are performing those assessments—for example, by helping them anticipate issues that could arise when physicians or nurses start working with an AI tool. They could also receive data from AI tool assessments performed by others and convene experts to consider whether the data support a recommendation that a particular AI tool be more widely used (*e.g.*, by including it in a clinical practice guideline), or that it be used or not used in a specific way.

⁴Mello MM, Rose S. Denial—artificial intelligence tools and health insurance coverage decisions. JAMA Health Forum (forthcoming on March 7, 2024; to be available at <https://jamanetwork.com/journals/jama-health-forum> or from the authors at mmello@law.stanford.edu).

Question. Rhode Island has two, unusually good Accountable Care Organizations (ACOs). In what ways could artificial intelligence be used to support the ACO program?

Answer. ACOs become eligible for Shared Savings by meeting certain quality measures. For example, several measures pertain to preventive care screening rates, such as mammography and fall risk screening. For some of these quality measures, AI tools could improve ACOs' ability to identify and reach patients in the relevant group for a particular measure. For instance, they could improve a hospital's ability to predict which patients are at highest risk of readmission within 30 days of discharge. AI tools could also help identify other quality measures that could be added to CMS's list (for example, by identifying aspects of care that predict better patient outcomes).

PREPARED STATEMENT OF ZIAD OBERMEYER, M.D., ASSOCIATE PROFESSOR AND BLUE CROSS OF CALIFORNIA DISTINGUISHED PROFESSOR, UNIVERSITY OF CALIFORNIA, BERKELEY

AI will transform medicine and the health-care system—for better or for worse, depending on how it is built and applied.

Thank you for the opportunity to address the committee. I am a professor and researcher at Berkeley, but my work on AI is grounded in another part of my identity, as a practicing emergency physician. Seeing patients, from academic hospitals in Boston to Tséhootsooi Medical Center in Fort Defiance, AZ, has given me a window into the miracles of modern medicine—awe-inspiring innovations in the diagnosis and treatment of disease, that have extended and improved life for millions.

Getting those miraculous tests and treatments to the right patient at the right time, though, is difficult. It requires processing enormous amounts of information, much of it imperfect and uncertain; errors have life-and-death consequences. Senators, I cannot imagine what it is like to do your job, but my guess is that you face situations like that every day. Throughout my 10 years of practicing medicine, I have agonized over missed diagnoses, futile treatments, unnecessary tests, and more. The collective weight of these errors, in my view, is a major driver of the dual crisis in our health-care system: suboptimal outcomes at very high cost.¹ AI holds tremendous promise as a solution to both problems.

By helping doctors and others in the health-care system make better decisions, I believe AI can both improve health and reduce costs—a rare combination. Let me share a few concrete examples drawn from my own work, to convey why I am so optimistic.

In the U.S. alone, 300,000 people experience sudden cardiac death every year.² What makes these events so tragic is that many of them are preventable: had we known a patient was at high risk, we would have implanted a defibrillator in her heart, to terminate the potential arrhythmias that cause sudden death, and save her life. Unfortunately, we are very bad at knowing who is at high risk. It's not just that we miss 300,000 opportunities every year to implant defibrillators and prevent those deaths. Even when we do implant defibrillators, we often do so in the wrong patients: up to one-third of patients end up never needing their defibrillator, meaning we increased their risk of complications and wasted resources without ever delivering a lifesaving shock.³

Fifteen years ago, when I was in medical school, I was stunned by these numbers. Today, working with colleagues in the U.S. and Sweden, we have trained an AI system to predict the risk of sudden cardiac death using just the waveform of a patient's electrocardiogram (ECG). It performs far better than our current prediction technologies, based largely on human judgment. This means we have the potential

¹While economists and policymakers have traditionally focused on the role of misaligned financial incentives, a large and growing body of recent research indicates much of our health-care system's inefficiency has its roots in human error. See Baicker, K., Mullainathan, S. and Schwartzstein, J., 2015. Behavioral hazard in health insurance. *The Quarterly Journal of Economics*, 130(4), pp. 1623–1667.

²Huikuri, H.V., Castellanos, A. and Myerburg, R.J., 2001. Sudden death due to cardiac arrhythmias. *New England Journal of Medicine*, 345(20), pp. 1473–1482.

³Moss, A.J., Greenberg, H., Case, R.B., Zareba, W., Hall, W.J., Brown, M.W., Daubert, J.P., McNitt, S., Andrews, M.L. and Elkin, A.D., 2004. Long-term clinical course of patients after termination of ventricular tachyarrhythmia by an implanted defibrillator. *Circulation*, 110(25), pp. 3760–3765.

to both save more lives *and* reduce waste, by ensuring that precious defibrillators are implanted in the right patients. It's rare to have an opportunity to both improve quality and reduce cost; normally we must choose. AI is a transformative new way for us to sidestep this dilemma entirely, and rebuild our health-care system on a foundation of data-driven decision-making.

This principle—better human decisions through AI-driven predictions—extends far beyond sudden cardiac death. We've found similar opportunities to improve quality and reduce cost in settings ranging from invasive testing for heart attack⁴ to mammograms for breast cancer prevention.⁵ AI can also help diagnose social vulnerability: we have promising early results showing AI can find subtle signs of interpersonal violence in x-rays. This means physicians can help recognize victims of violence when they come to seek help in the ER—instead of missing those opportunities, as happens all too often today—and social services can connect them to the resources they need.⁶ AI is also starting to drive innovation in the science of medicine, for example, by discovering entirely new classes of antibiotics in drug libraries that were passed over for decades by human researchers.⁷

Despite my great optimism, I worry that without concerted effort from researchers, the private sector, and government, AI may be on a path to do more harm than good in health care. So I'd also like to provide an example of how AI can go wrong. Working on this problem has taught me a lot about we can work together to ensure that AI systems are safe.

Five years ago, my colleagues and I uncovered evidence that a family of poorly designed AI algorithms, built and used in both public and private sectors, contained large-scale racial bias.⁸ These algorithms had a laudable goal: to identify patients at high risk of future health problems—exacerbations of chronic conditions like heart failure, diabetes, etc. The AI's predictions are used by health systems around the world to decide who gets access to extra help, in the form of “care management” programs. In theory, this is a great use of AI, because these programs are a win-win: high-risk patients get the help they need to manage chronic conditions, reducing future flare-ups and complications; and the health-care system saves the money it would have spent on the resulting ER visits and hospitalizations.

Unfortunately, a subtle-seeming choice in the AI's design caused untold harm: a gap between what the algorithms were *supposed* to predict (health-care needs) and what they *actually* predicted (health-care costs). The AI's goal was to identify patients with high future health needs. But AI is extremely literal—it predicts a specific variable, in a specific dataset—and there is no variable available called “future health needs.” So instead, the AI developers chose to predict a proxy variable that *is* present in health datasets: future health-care *costs*. Spending on health care seems like a reasonable proxy for health needs. After all, sick people generate health costs. But because of discrimination and barriers to access, underserved patients who need health care often don't get it. This means Black patients—and also poorer patients, rural patients, less-educated patients, and all those who face barriers to accessing health care when they need it—get less spent on their health care than their better-served counterparts, even though they have the same underlying health conditions.⁹ Low costs do not necessarily mean low needs.

Tragically, the AI ignored these simple facts. It predicted—accurately—that Black patients would generate lower costs, and thus deprioritized them for access to help with their health. The result was racial bias that affected important decisions for hundreds of millions of patients every year. Senator Wyden, I was heartened by the letters that you and Senator Booker sent to executives at major insurance companies in the wake of that study—I believe that had a great impact. Unfortunately,

⁴Mullainathan, S. and Obermeyer, Z., 2022. Diagnosing physician error: A machine learning approach to low-value health care. *The Quarterly Journal of Economics*, 137(2), pp. 679–727.

⁵Daysal, N.M., Mullainathan, S., Obermeyer, Z., Sarkar, S.K. and Trandafir, M., 2022. An Economic Approach to Machine Learning in Health Policy. CEBI Working Paper. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4305806.

⁶Williams, B., Oto, A., Ludwig, J., Graber, R., Obermeyer, Z. and Mullainathan, S., 2023. Making the invisible epidemic visible. <https://www.brookings.edu/articles/making-the-invisible-epidemic-visible/>.

⁷Stokes J.M., et al. A deep learning approach to antibiotic discovery. *Cell*. 2020 February 20;180(4):688–702.

⁸Obermeyer, Z., Powers, B., Vogeli, C. and Mullainathan, S., 2019. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), pp. 447–453.

⁹While our work focused on demonstrating one specific bias—Black versus White patients—this is not an issue of race alone: any populations with a wedge between the care they need and the care they get will be similarly affected.

many of the biased algorithms we studied remain in use today. And similar dynamics were highlighted by a recent investigation of AI products used to deny claims:¹⁰ in all these cases, AI learns from historical data, with all its biases and inequities, and encodes those past practices in policy. So those underserved patients whose claims have been denied by humans in our past datasets—often for unjust reasons—will have their claims denied by AI at scale, forever, unless we can realign AI with our society’s goals.

Fortunately, there are a number of specific things that programs under this committee’s jurisdiction can do to ensure that AI produces the social value we all want. I believe that Medicare, Medicaid, CHIP, and child welfare programs stand to realize enormous benefits from AI: well-designed products can both improve the quality of services and reduce their cost. As a result, these programs should be willing to pay for AI—but they should not simply accept the flawed products that the market often produces. Rather, they should take advantage of their market power to articulate clear criteria for what they will pay for, and how much. I believe this will harness the tremendous innovative power of the market, and ensure it is pointed in the right direction.

Based on my research, as well as my work with Federal regulators and State Attorneys General, I believe that programs in the committee’s jurisdiction should explicitly evaluate AI algorithms for reimbursement on a small set of targeted criteria. AI developers must be transparent about the output of their algorithms. If an algorithm predicts health care costs, the developer should not be able to claim that it predicts “health risks” or “health needs”—unfortunately, many cost-predictors currently do exactly this. Algorithms’ outputs should be evaluated for accuracy in a completely independent dataset, both overall and in protected groups, in keeping with good machine learning practice.¹¹ I emphasize that this approach focuses on the *output* of the algorithm—the accuracy of its predictions on a transparently stated target—and thus does not require “opening the black box”: algorithms can be evaluated simply based on the predictions they produce. This avoids compromising trade secrets, and means purchasers (and similarly, regulators) do not need to evaluate the *inputs* of algorithms, or understand the many reasons why they might be biased—technical problems with a complex model, non-representative training data, use of an explicit race correction, etc. Instead, we can focus on one simple question: is the algorithm predicting what it’s supposed to predict, accurately and equitably?¹² Finally, AI products should be valued and reimbursed according to established principles from health economics and outcomes research. If an AI results in an earlier diagnosis of heart attack or breast cancer, for example, that generates value to patients in the form of life-years, and to the health-care system in the form of downstream costs avoided. The sooner public programs lay out what they are looking for, the sooner the market can deliver safe and effective AI products to solve the urgent problems they face.

I should note that my applied work has resulted in collaborations with a number of public and private entities, but the views I present are entirely my own, based on my experiences and research.

Many thanks again for this opportunity. I look forward to answering your questions.

¹⁰Ross, C. and Herman, B., 2023. Denied by AI. *STAT News*. March 13, 2023. <https://www.statnews.com/2023/03/13/medicare-advantage-plans-denial-artificial-intelligence/>.

¹¹Accessing data for such evaluations is a non-trivial problem, but there are emerging solutions. For example, a company I cofounded, Dandelion Health, offers a free public service for the evaluation of algorithm performance and equity. This service, which is philanthropically supported by the Gordon and Betty Moore Foundation and the SCAN Foundation, allows any AI developer to securely upload the algorithm to Dandelion’s computing environment. Dandelion will run the algorithm on its diverse national dataset and deliver back a report on the algorithm’s performance, both overall and across key geographic, racial, ethnic, age, gender, and socioeconomic groups. More details are at <https://dandelionhealth.ai/validation>.

¹²This is analogous to the process by which the FDA regulates information about drugs: pharmaceutical companies must be transparent about the primary outcome a drug is intended to improve, and the drug’s impact on that outcome is assessed in a rigorous randomized trial.

QUESTIONS SUBMITTED FOR THE RECORD TO ZIAD OBERMEYER, M.D.

QUESTIONS SUBMITTED BY HON. RON WYDEN

Question. Providers use algorithms to guide clinical care—making diagnoses, recommending treatments, and predicting outcomes. Unfortunately, some of these algorithms, including one that you investigated in 2019, have been built atop flawed data and faulty assumptions. These tools can perpetuate harmful biases and disparities in care—ultimately hurting patients. It’s another example where the profits are being privatized, but the risks are being socialized.

What should Congress consider to ensure that clinical AI tools are used to support an equitable health-care system so that everyone benefits?

There has been significant concern raised about racial and other biases in predictive algorithms. In 2022, I urged my home State of Oregon to stop using its AI screening tool because of reports that Black families were being disproportionately referred for mandatory child neglect investigations when the tool was used. I am glad that Oregon paused the use of their tool—making decisions about families is too important to use tools that aren’t ready for prime time and bake in bias.

What steps does Congress need to take to ensure that child welfare agencies are using AI and predictive algorithms in responsible and ethical ways, avoiding bias, and ensuring privacy protection for the personal data of children and families?

Answer. I believe Congress has two different kinds of opportunities to ensure that AI tools are developed in an effective and just way, for both health and child welfare.

The most powerful lever Congress, and the Senate Finance Committee in particular, can pull is incentives. Today, *AI developers are getting little guidance from payers or government agencies on which products will be reimbursed.* That uncertainty leads to under-investment in developing new products that benefit patients and society, and has also resulted in deeply flawed tools that perpetuate bias.

CMS can solve this problem by defining clear criteria, and accompanying payment rates, for health AI products that improve the cost-effectiveness of care. This will set the goalposts for industry innovation in a way that benefits the public, as well as CMS. I am very optimistic that AI can have meaningful health and economic impacts, by detecting and treating disease early in a way that helps providers change the disease trajectory. Working across a range of high-priority problems—e.g., dialysis, end-stage heart failure, frailty and osteoporosis—standard cost-effectiveness analysis can quantify the value of AI at the patient level. This can be used to decide whether and how much CMS should pay health providers to run the AI on a single patient.

Much like an advance market commitment,¹ CMS could then announce the reimbursement level for providers to run a specific AI tool on a patient. The price would be based on cost-effectiveness; it could and should be higher than marginal cost of running the algorithm: the goal is to catalyze investment. The price would then taper off on a set schedule, then could transition to value-based reimbursement for the outcome we want (e.g., preventing dialysis).

Importantly, *reimbursement should be conditional on the AI meeting performance and equity (race/ethnicity, geography, et cetera) targets*, demonstrated in an independent third-party sample the developer did not use to train the AI. The level of reimbursement could even be tied to the level of AI performance, to spur competition and avoid creating a monopoly.

The second lever Congress can pull is direct regulatory oversight of AI tools, by requiring developers to adhere to two simple requirements before tools are used for patient care.

First, AI developers must be *transparent about the output of their algorithms*. If an algorithm predicts health-care costs, the developer should not be able to claim that it predicts “health risks” or “health needs.” I believe this will help surface what you rightly referred to as the “flawed data and faulty assumptions” that underlie many current tools.

¹ Kremer, M., Levin, J., and Snyder, C.M., 2020, May. Advance market commitments: insights from theory and experience. In *AEA Papers and Proceedings* (Vol. 110, pp. 269–273).

Second, *the outputs of AI tools should be evaluated for accuracy in a completely independent dataset*, both overall and in protected groups, in keeping with good machine learning practice. By focusing on the *output* of the algorithm—the accuracy of its predictions on a transparently stated target—regulators and customers do not need to “open the black box”: algorithms can be evaluated simply based on the predictions they produce. This avoids compromising trade secrets, and means purchasers (and similarly, regulators) do not need to evaluate the *inputs* of algorithms, or understand the many reasons why they might be biased—technical problems with a complex model, nonrepresentative training data, use of an explicit race correction, et cetera. Instead, we can focus on one simple question: is the algorithm predicting what it’s supposed to predict, accurately and equitably?

I should note that accessing data for such evaluations is a non-trivial problem. Here too, there are steps that Congress and agencies under its jurisdiction can take. First, health and child welfare *agencies should make their data available to researchers*, to facilitate independent analyses of algorithms.

Second, Congress should support nascent efforts to create *public-private partnerships that make available high-quality, diverse datasets for AI training and validation*. For example, the FDA is investigating a “national quality assurance lab” model where AI developers could securely upload their algorithms to independent dataset to objectively test their accuracy and equity.²

QUESTIONS SUBMITTED BY HON. SHELDON WHITEHOUSE

Question. Rhode Island has several health-care enterprises that we are very proud of, including the State-wide Rhode Island All-Payer Claims Database (RI APCD).

What role can AI play with regard to maximizing the utility of a robust, All-Payer Claims Database?

Are you working in the development of AI with some of the medical specialty organizations (orthopedics, cardiologists, et cetera)? Are specialty organizations a useful place for benchmarking, approval, or accrediting best AI practices? And if they are, do you have any good examples of a medical specialty association that is being particularly forward and helpful at looking for the best uses of AI within the specialty?

Rhode Island has two, unusually good Accountable Care Organizations (ACOs).

In what ways could artificial intelligence be used to support the ACO program?

Answer. I believe the single most important thing that Rhode Island can do to maximize the value of its all-payer claims database is to *streamline the access and approval process to make the data available more broadly outside of Rhode Island*. I have seen firsthand the amazing work that has been produced based on these data by researchers at Brown. Facilitating the process for new researchers to access the data would dramatically scale up their impact, far beyond Rhode Island. There are technological and legal frameworks that keep data secure and protect privacy: storing the data on modern cloud computing platforms is very secure; contractual and legal safeguards do not need to be onerous to be effective.

I believe AI tools are particularly suited to the incentive environment of ACOs. By providing doctors and policymakers with accurate predictions, AI can enhance efforts that prevent disease at low cost, rather than treat it after it gets out of control. This applies to diagnosing heart attack,³ cancer screening tools like mammograms,⁴ and even social vulnerability: we have promising early results showing AI

²As a concrete example of such an effort, a company I cofounded, Dandelion Health, offers a free public service for the evaluation of algorithm performance and equity. This service, which is philanthropically supported by the Gordon and Betty Moore Foundation and the SCAN Foundation, allows any AI developer to securely upload the algorithm to Dandelion’s computing environment. Dandelion will run the algorithm on its diverse national dataset and deliver back a report on the algorithm’s performance, both overall and across key geographic, racial, ethnic, age, gender, and socioeconomic groups. More details are at <https://dandelionhealth.ai/validation>.

³Mullainathan, S. and Obermeyer, Z., 2022. Diagnosing physician error: A machine learning approach to low-value health care. *The Quarterly Journal of Economics*, 137(2), pp. 679–727.

⁴Daysal, N.M., Mullainathan, S., Obermeyer, Z., Sarkar, S.K. and Trandafir, M., 2022. An Economic Approach to Machine Learning in Health Policy. CEBI Working Paper. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4305806.

can find subtle signs of interpersonal violence in x-rays. This means physicians can help recognize victims of violence when they come to seek help in the ER—instead of missing those opportunities, as happens all too often today—and social services can connect them to the resources they need.⁵

This means we have the potential to both save more lives *and* reduce waste, by ensuring that precious resources are allocated to the right patients. It's rare to have an opportunity to both improve quality and reduce cost; normally we must choose. AI is a transformative new way for us to sidestep this dilemma entirely, and rebuild our health care system on a foundation of data-driven decision-making.

I am not working with any medical specialty organizations. I believe those organization could have a productive role in convening AI developers and urging them to follow best practices (*e.g.*, as noted in my responses to Senator Wyden—specificity and independent evaluation) as they develop their tools.

PREPARED STATEMENT OF MARK SENDAK, M.D., MPP,
CO-LEAD, HEALTH AI PARTNERSHIP

Chairman Wyden, Ranking Member Crapo, and members of the committee, my name is Mark Sendak, and I appreciate the opportunity to serve on the panel today and offer my testimony. I must note that any views expressed today and in my written testimony are my own and may not necessarily reflect those of my employer or the institutions that participate in a multi-stakeholder partnership of which I am proud to serve in a leadership role.

I serve as the Population Health and Data Science Lead at the Duke Institute for Health Innovation (DIHI for short) and the co-lead of Health AI Partnership. I studied mathematics and health policy before completing medical training and have been developing and implementing AI technologies for over a decade.

Since DIHI's founding in 2013, our team has led the development and responsible implementation of over 20 AI technologies for clinical care.¹ We were the first in the U.S. to implement a deep learning model in routine clinical care.² We were the first to implement Model Facts labels (similar to nutrition facts labels) for AI tools that laid the groundwork for ONC's final rule on algorithm transparency.^{3,4} We have incubated four companies to commercialize AI products built at Duke and we help health-care organizations across the country validate these technologies.

Our team has demonstrated the benefits of AI in health care. Duke dramatically improved the quality of sepsis care using our Sepsis Watch system.⁵ Duke proactively manages chronic diseases in Medicare patients by using AI to identify patients at risk of complications.⁶ Duke has now extended this chronic disease management approach to all patients.

But my comments today will not focus on the amazing work I've been a part of at Duke. Today, I'm speaking with you primarily as the co-lead of Health AI Partnership.

In 2018, a mentor asked me "how do we get AI out of the ivory tower?" At that time, my experience with AI at Duke was unimaginable to people outside a few exceptional islands of excellence. And there was minimal building of infrastructure to advance the use of AI in low-resource settings.

⁵Williams, B., Oto, A., Ludwig, J., Graber, R., Obermeyer, Z. and Mullainathan, S., 2023. Making the invisible epidemic visible. <https://www.brookings.edu/articles/making-the-invisible-epidemic-visible/>.

¹Sandhu S, Sendak MP, Ratliff W, Knechtle W, Fulkerson WJ, Balu S. Accelerating health system innovation: principles and practices from the Duke Institute for Health Innovation. *Patterns*. 2023;4(4):100710.

²Sendak MP, Ratliff W, Sarro D, Alderton E, Futoma J, Gao M, et al. Real-World Integration of a Sepsis Deep Learning Technology Into Routine Clinical Care: Implementation Study. *JMIR medical informatics*. 2020 Jul 15;8(7):e15182.

³Sendak MP, Gao M, Brajer N, Balu S. Presenting machine learning model information to clinical end users with model facts labels. *npj Digital Medicine*. 2020 Mar 15;3(41):1–4.

⁴<https://www.healthit.gov/topic/laws-regulation-and-policy/health-data-technology-and-interoperability-certification-program>.

⁵<https://www.wsj.com/articles/how-hospitals-are-using-ai-to-save-lives-11649610000>.

⁶Sendak MP, Balu S, Schulman KA. Barriers to Achieving Economies of Scale in Analysis of EHR Data. A Cautionary Tale. *Applied Clinical Informatics*. 2017 Aug 9;8(3):826–31. Available from: <https://www.thieme-connect.com/products/ejournals/abstract/10.4338/ACI-2017-03-CR-0046>.

In 2021, I helped launch Health AI Partnership to advance the safe, effective, and equitable use of AI in all health-care organizations. We exist to get AI out of the ivory tower.

The Senate Finance Committee can take concrete action to advance accountability, equity, privacy, and transparency in the use of AI in health care. The Medicare program ensures the delivery of high-quality care for beneficiaries through conditions of participation and other mechanisms. There is a unique opportunity for this committee to strengthen Medicare controls on the use of AI and to facilitate investments in technical assistance, technical infrastructure, and training.

First, we can address guard rails.

Through Health AI Partnership, we work with 20 organizations across the U.S. to surface and disseminate AI best practices. We interview leaders and run case-based workshops on complex topics. We develop practical resources for health-care leaders asking basic questions: how do I evaluate different externally built AI products? How do I navigate the new FDA clinical decision support guidance? How do I assess the potential future impact of this AI product on health inequities? How do I align organizational processes with the White House Blueprint for an AI Bill of Rights?

Health AI Partnership resources and programs provide guardrails for high-resource organizations that are rapidly accelerating their use of AI. Adoption of these guard rails by hospitals could be required for Medicare program participation. But guard rails only serve the few organizations that are already on the AI adoption highway.

We must also address the more critical need for roads, onramps, and bridges—the core infrastructure investments needed to ensure that all people in the U.S. benefit from AI in health care.

Most health-care organizations in the U.S. need an on ramp to the AI adoption highway. They are struggling with clinician burnout. They face razor thin or negative margins. They are entirely dependent on external EHR vendors for technology expertise and assistance. Simply put, they do not have the resources, personnel, or technical infrastructure to embrace guardrails for the AI adoption highway.

Core infrastructure investments are needed for technical assistance, technology infrastructure, and training:

- **Technical Assistance:** A national hub-and-spoke network is needed to diffuse AI expertise beyond centers of expertise to low-resource settings. Hubs can provide operational and technical support to sites implementing AI, similar to existing programs that extend specialist expertise to low-resource settings.^{7,8}
- **Technology Infrastructure:** Technology infrastructure that is distinct from EHRs is required to facilitate the efficient evaluation and clinical integration of AI tools. This infrastructure allows sites to test many AI products simultaneously with ongoing monitoring. Without addressing this capital investment, a market failure will continue preventing AI developers from efficiently commercializing products.⁹
- **Training programs:** Broadly accessible programs targeting clinical, technical, and operational leaders are urgently needed to equip the health-care workforce with the foundational knowledge required to locally govern AI. Health care will increasingly become AI-enabled and local AI governance will be a core competency.

Congress has tackled this type of challenge before. Fifteen years ago, Congress enabled the broad adoption of EHRs through funding technical assistance programs and technology infrastructure investments.¹⁰ That funding supported the purchase of EHRs along with 62 regional extension centers to support EHR implementations in low-resource settings. While EHRs are far from perfect, Federal programs did

⁷<https://projectecho.unm.edu/>.

⁸<https://dason.medicine.duke.edu/>.

⁹<https://www.statnews.com/2022/05/24/market-failure-preventing-efficient-diffusion-health-care-ai-software/>.

¹⁰Lynch K, Kendall M, Shanks K, Haque A, Jones E, Wanis MG, et al. The Health IT Regional Extension Center Program: Evolution and Lessons for Health Care Transformation. *Health Serv Res.* 2014;49(1pt2):421–37.

successfully diffuse the technology across the country. We need similarly bold action now.

Enacting Medicare controls and infrastructure investments to advance the safe, accountable, equitable, private, and transparent use of AI in health care will take time and be iterative. I look forward to continuing to share learnings from Health AI Partnership and am eager to support future work conducted by the Senate Finance Committee.

Thank you, again, for this opportunity, and I look forward to answering your questions.

QUESTIONS SUBMITTED FOR THE RECORD TO MARK SENDAK, M.D., MPP

QUESTIONS SUBMITTED BY HON. RON WYDEN

Question. For years I've been insisting that companies need to step up and assess the impacts of the AI systems they develop, which is what my Algorithmic Accountability Act would require. Unfortunately, many companies have done far too little to make sure their automated black-box systems really work and do not amplify bias and discrimination. Right now, neither the public nor the government knows when, whether, or how AI systems are being used to make critical decisions about Americans' homes, finances, and jobs. Especially in health care, where AI systems can directly impact the health of Americans, accountability and transparency are absolutely critical.

What rules should companies building AI tools for health care abide by? What should we expect of providers who use AI tools?

Answer. Companies that build and implement AI tools for health care should:

- Transparently share information about the AI tool with health care delivery organizations (HDO) and the public. Information about the AI tool to curate and transparently report can align with the Health Data, Technology, and Interoperability (HTI-1) final rule: <https://www.healthit.gov/topic/laws-regulation-and-policy/health-data-technology-and-interoperability-certification-program>. Information should include how the tool was built, the intended use of the tool, how the tool should be integrated into clinical workflows, information about the data used to train the model, as well procedures to assess the potential impact of the AI tool on health inequity. The Health Equity Across the AI Lifecycle (HEAAL) framework describes potential procedures to complete to assess health equity impacts of AI tools: <https://www.medrxiv.org/content/10.1101/2023.10.16.23297076v5>.
- Provide an interface for clinicians, patients, HDO leaders, and other affected stakeholders to report adverse events related to use of the AI tool. The company should provide periodic public disclosure of adverse events as well as any action taken to improve AI product performance and use. If the company becomes aware of an issue that can cause harm to patients, the company should immediately report the issue to clinicians, patients, and HDOs that use the tool.
- Provide technical assistance, training, and technical infrastructure capabilities to support HDOs rigorously evaluate and monitor the AI product.

Specific HDO activities that the company should support include:

- Assess the performance of the AI tool on historical data within the implementation context. The company should provide guidance on performance thresholds that support progressing to a "silent trial."
- Conduct a "silent trial" in which the AI tool is run prospectively and not used clinically within the implementation context. The company should provide guidance on performance thresholds that support clinical integration.
- Continuously monitor performance of an AI tool that is clinically integrated. The company should provide guidance on performance thresholds that support continued clinical use.

HDOs that use AI tools should:

- Adopt best practices for AI product lifecycle management, such as those set out by Health AI Partnership: <https://healthaipartnership.org>. A third-party

entity can assess HDO implementation and adoption of the best practices to ensure that HDOs rapidly strengthen internal controls. Market incentives should reward HDOs that successfully adopt AI product lifecycle management best practices.

- Establish a risk-based formal review and approval process for clinical and operational AI product use. High-risk use cases should be subject to additional scrutiny and controls.
- Conduct multiple forms of AI product assessments, including:
 - Assess the performance of the AI tool on historical data within the implementation context. The HDO should publicly report a subset of the performance measure data.
 - Conduct a “silent trial” in which the AI tool is run prospectively and not used clinically within the implementation context. The HDO should publicly report a subset of the performance measure data.
 - Continuously monitor performance of an AI tool that is integrated into clinical care or operations. The HDO should publicly report a subset of the performance measure data.
- There must be market incentives to support HDO procurement and maintenance of technical infrastructure required to conduct the AI product testing described today. These resources should be separate from reimbursement for AI product adoption and should be considered infrastructure investments. For more on the need for these investments, please see: <https://www.statnews.com/2022/05/24/market-failure-preventing-efficient-diffusion-health-care-ai-software/>.
- Publicly report significant changes in AI product performance that are observed after clinical integration of the tool. This information should be shared with a third-party entity that curates and aggregates information from HDOs across the country.
- Periodically audit AI tools used in clinical care and operations. The HDO can work with a third-party entity to conduct the audit and there should be market incentives to support completion of the algorithmic audit. A subset of the audit results should be publicly reported.
- Publicly disclose high-risk AI tools that are used in clinical care or operations. Information about the tool should be presented in a format that is useful to patients and addresses questions and concerns that patients have about the AI tool. Health care delivery organizations can adopt a risk-based approach to focus disclosure efforts on AI use cases that are at high-risk of worsening health inequities, target high-risk conditions, and play a significant role shaping access to health care and delivery of health care.

Question. There has been significant concern raised about racial and other biases in predictive algorithms. In 2022, I urged my home State of Oregon to stop using its AI screening tool because of reports that Black families were being disproportionately referred for mandatory child neglect investigations when the tool was used. I am glad that Oregon paused the use of their tool. Making decisions about families is too important to use tools that aren’t ready for prime time and bake in bias.

What steps does Congress need to take to ensure that child welfare agencies are using AI and predictive algorithms in responsible and ethical ways, avoiding bias, and ensuring privacy protection for the personal data of children and families?

Answer. Congress should:

- Require developers of AI products used by child welfare agencies to transparently report information about the AI product. Information about the AI tool to curate and transparently report can align with the Health Data, Technology, and Interoperability (HTI-1) final rule: <https://www.healthit.gov/topic/laws-regulation-and-policy/health-data-technology-and-interoperability-certification-program>. Information should include how the tool was built, the intended use of the tool, how the tool should be integrated into clinical workflows, information about the data used to train the model, as well procedures to assess the potential impact of the AI tool on health inequity. The Health Equity Across the AI Lifecycle (HEAAL) framework describes potential procedures to complete to assess health equity impacts of AI tools: <https://www.medrxiv.org/content/10.1101/2023.10.16.23297076v5>.

- Require developers of AI products used by child welfare agencies to provide an interface for front-line service providers and affected members of the public to report adverse events related to use of the AI tool. The company should provide periodic public disclosure of adverse events as well as any action taken to improve AI product performance and use. If the company becomes aware of an issue that can cause harm to members of the public, the company should immediately report the issue to child welfare agencies that use the tool.
- Require child welfare agencies to publicly disclose high-risk AI tools that are used in service delivery or operations. Information about the tool should be presented in a format that is useful to the public and addresses questions and concerns that members of the public have about the AI tool.
- Require child welfare agencies to conduct multiple forms of AI product assessments, including:
 - Assess the performance of the AI tool on historical data within the implementation context. The child welfare agency should publicly report a subset of the performance measure data.
 - Conduct a “silent trial” in which the AI tool is run prospectively and not used clinically within the implementation context. The child welfare agency should publicly report a subset of the performance measure data.
 - Continuously monitor performance of an AI tool that is integrated into clinical care or operations. The child welfare agency should publicly report a subset of the performance measure data.
- Require child welfare agencies to periodically audit AI tools used in service delivery and operations. The child welfare agency can work with a third-party entity to conduct the audit and there should be market incentives to support completion of the algorithmic audit. A subset of the audit results should be publicly reported.

Beyond child welfare, Congress can require that organizations that provide health and human services conduct health equity impact assessments for high-risk AI applications. Health AI Partnership convened leaders from HDOs across the United States to present real-world cases of AI being evaluated for clinical use and surfaced a set of procedures across the AI product lifecycle to mitigate the risk of worsening health inequities. The resultant framework is called Health Equity Across the AI Lifecycle and it can be found here: <https://www.medrxiv.org/content/10.1101/2023.10.16.23297076v5>. Congress can require that for certain high-risk applications, organizations that provide health and human services conduct the procedures and publicly report a subset of the results. There are significant costs associated with conducting the health equity impact assessments and Congress should provide market incentives for organizations to conduct the assessments.

QUESTIONS SUBMITTED BY HON. CHUCK GRASSLEY

Question. Do the Centers for Medicare and Medicaid Services (CMS) have the capability to determine an accurate and fair reimbursement level for artificial intelligence technologies, so we don’t waste taxpayer money?

Answer. The capability to determine accurate and fair reimbursement for AI requires both personnel and high-quality, granular, real-world performance data. CMS should prioritize reimbursement for AI products that generate value in real-world HDO implementation contexts. Given that AI product performance can vary across HDO contexts, HDOs should regularly report AI product performance data to CMS. If the performance of an AI product changes within a HDO implementation context and the product no longer generates value, CMS should adjust payment rates to the HDO accordingly.

CMS should convene stakeholders across clinical, technical, and business domains to assess the value creation and value capture of AI products being considered for clinical and operational use. This interdisciplinary advisory group could complement and augment the Medicare Payment Advisory Committee (MedPAC), which advises the U.S. Congress on issues affecting the Medicare program (<https://www.medpac.gov/>). The advisory group should include community and patient representatives who are affected by AI use in health care.

CMS, along with other HHS agencies, will greatly benefit from the National AI Talent Surge announced in the executive order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. There is an urgent need to increase the hiring and retention of clinical and technical experts with practical experience developing, implementing, and maintaining AI products used in clinical care. Congress should continue to fund and expand efforts that expand CMS capabilities in AI.

Given that AI will become increasingly integral to health-care delivery, CMS needs to develop a business model for AI reimbursement that accounts for the total product lifecycle of AI tools. There are significant costs associated with the development and maintenance of technical infrastructure required to ensure that AI advances safe, effective, and equitable care. These costs are fixed and often represent capital investments that are amortized across many AI applications.

Question. We spent about \$4.5 trillion in health care last year. A key area of waste in our health-care system is medical errors, failure of care delivery, and over-treatment. Some estimates suggest we waste \$205 billion to \$425 billion each year due to medical errors and waste.

What potential does deploying artificial intelligence to review clinical decisions have in reducing medical errors and waste? Are you aware of existing practices and can we scale them?

Answer. There are many applications of clinical decision support (CDS) that target acute conditions to both improve detection and coordinate appropriate and timely intervention. These CDS technologies often use AI to predict complications before they occur and serve as a safety net to catch clinical deterioration that may be missed by burdened front-line clinicians. An example of this type of system at Duke Health is Sepsis Watch, which monitors the care provided to all adult patients who present to the emergency department. Sepsis Watch is used by a centralized team of nurses to identify patients at risk of deterioration who are in need of evidence-based interventions. The nurses support front-line clinicians to complete relevant treatment interventions. In order to scale a tool like this, there needs to be reimbursement for health care delivery organizations (HDOs) to run the tool, the technical infrastructure required to integrate the AI product into the electronic health record and to monitor performance of the AI product, and the personnel required to utilize the AI product to support front-line clinicians.

In a fee-for-service payment environment, medical errors and medical complications can serve as a source of revenue for HDOs. Reimbursement models need to be established to fund the integration and diffusion of AI products that prevent errors and medical complications.

Clinician burnout remains alarmingly high and provider burden is a main driver of medical errors and medical complications. Large language models (LLMs) provide an opportunity to leverage AI to relieve the documentation and administrative burden placed on front-line clinicians. For example, AI scribe technologies are being adopted by HDOs to draft visit notes from recordings of clinic encounters. These technologies can be validated and scaled through reimbursement mechanisms that facilitate adoption.

QUESTIONS SUBMITTED BY HON. JOHN CORNYN

Question. The number of Americans living with chronic diseases is expected to increase as our population ages. The NIH estimates that by 2050 there will be nearly 143 million Americans over 50 living with at least one chronic disease. This is up from about 72 million Americans in 2020. Such a dramatic increase will put even more pressure on the Medicare system. AI has the possibility of helping providers more efficiently identify patients at risk of chronic diseases. This could include patients at risk of developing type 2 diabetes or other preventable diseases.

In your testimony you mentioned some of the work Duke is doing to proactively manage chronic diseases in Medicare patients.

Can you talk more about some of the successes you have seen in using AI to identify and prevent chronic diseases?

Answer. Duke Health develops and implements numerous AI tools to predict chronic disease progression to intervene and prevent downstream, costly complications. We have done this through a combination of technology and workflow innova-

tion, developing a process called population rounding where an interdisciplinary team reviews high-risk patients and sends tailored recommendations to primary care physicians. For example, we implemented a chronic kidney disease algorithm to facilitate interventions that prevent kidney failure, a peripheral artery disease algorithm to prevent lower limb amputations, a cardiovascular risk algorithm to prevent heart attacks, and a mortality algorithm to facilitate advance care planning and community-based palliative care. These AI products were all implemented within a Medicare Shared Savings Program accountable care organization, which creates financial incentives for prevention of chronic disease progression.

Duke Health is a national leader in the development and scaling of AI governance, which supports the use of AI for chronic disease management. Clinical and technical teams across the health enterprise, school of medicine, and university campus collaborate closely to operate the Algorithmic-Based Clinical Decision Support oversight. The effort involves leadership from DIHI, Duke AI Health, and Duke Health Technology Solutions. Duke AI Health connects, strengthens, amplifies, and grows multiple streams of theoretical and applied research on AI and machine learning at Duke University in order to answer the most urgent and difficult challenges in medicine and population health.

Question. What are the opportunities you see for using AI to lower Medicare costs for seniors?

Answer. Opportunities to use AI to lower Medicare costs include:

- Leverage AI to tailor disease monitoring and surveillance to individual patient risk to move beyond coarse screening protocols. This could reduce the amount of unnecessary and low-value testing.
- Leverage AI to monitor patients at home through remote sensors. Enhanced remote monitoring can detect physiologic deterioration early to prompt telemedicine interventions and prevent progression that would result in acute care visits and hospital admissions.
- Leverage AI to route patients to appropriate care setting based on data gathered through remote sensors as well as through LLM interactions. This can help ensure that patients who need acute care have their care rapidly escalated and can also help prevent unnecessary emergency department visits and hospital admissions.
- Leverage AI to optimize medical prescriptions to lowest cost therapeutic options that are effective for patients.
- Leverage AI to assess risk of surgical complications to optimize preoperative conditioning and to route patient to most cost-effective operative setting. This can help route patients at low risk of complication to ambulatory surgery centers and route patients at high risk of complication to inpatient surgery units.
- Patients would need to consent to use of data for these purposes and for interactions with LLMs.

QUESTION SUBMITTED BY HON. JOHN THUNE

Question. It is important to me as the Senate considers different approaches to regulating AI, that we don't create duplicative regulatory pathways. For example, the FDA already has a pathway in place to regulate types of AI, such as software as a medical device.

At the same time, the regulatory agencies also need to make continual improvements to their processes to address updates to the technology and there is a need for technical support and guidance to navigate these processes.

In your written testimony you highlight the need for a hub-and-spoke network to diffuse AI expertise and technical support. My bill supports this approach as well.

Can you expand on what the hub-and-spoke model might look like and why it is important?

Answer. The performance of an AI product across health care delivery organization (HDO) contexts can be highly variable and dynamic and depend on multiple factors including information technology (IT) systems, clinical workflows, patient populations, and local epidemiology. Even if an HDO confirms that an AI product performs well, the tool must be continuously monitored to ensure that performance

remains consistent. For these reasons, it is critical for HDOs to be able to locally validate and monitor AI products used in clinical care and operations. This requires the rapid diffusion of capabilities and technology.

A hub-and-spoke model would first formalize a network of “hub” sites that are mature in their use of AI products, agree upon best practices across the AI product lifecycle, share learnings to contribute to best practices for emerging challenges, and are committed to providing technical assistance and training to “spoke” sites that are early in their journey of AI adoption. A hub-and-spoke model is highly scalable and modular, facilitating expansion to new geographic regions, expansion to new medical conditions, and expansion to new generations of technology. A hub-and-spoke model can rapidly build capacity across diverse HDOs and rapidly stimulate economic growth to support the safe, effective, and equitable implementation of AI.

A hub-and-spoke model is highly redundant, ensuring robust health-care delivery and robust monitoring of AI product performance across contexts. For example, if a single “hub” is taken offline for technical or operational reasons, other “hubs” can step in to provide remote expertise and technical assistance to “spoke” sites trying to use AI. By facilitating redundancy, the hub-and-spoke model also facilitates rapid testing and experimentation of differing approaches to AI product lifecycle management. Given the speed of technology progress, providing “hub” sites flexibility in how they support “spoke” sites allows for the rapid emergence and diffusion of best practices.

A hub-and-spoke model spurs competition and innovation in the AI product development ecosystem by providing multiple avenues of market entry. By decentralizing control over market access across the network of hub-and-spoke sites, AI product developers can more easily develop products that target markets that may initially be quite niche.

A hub-and-spoke model can surface best practices for AI product use much more rapidly than centralized AI product testing and validation. For example, “hub” sites can individually demonstrate success for a given AI product use case to facilitate diffusion across the network. This allows many different sites to operate in parallel to test and diffuse best practices.

Lastly, while centralized AI product testing may seem more cost-effective and efficient, it does not scale. The number of HDOs that need to locally validate and monitor AI products is in the thousands and the number of AI products coming to market is growing exponentially. Tight, central control over market access for AI products will significantly slow and stifle innovation and will limit the potential impact of AI on health care.

QUESTIONS SUBMITTED BY HON. SHELDON WHITEHOUSE

Question. Rhode Island has several health-care enterprises that we are very proud of, including the State-wide Rhode Island All-Payer Claims Database (RI APCD).

What role can AI play with regard to maximizing the utility of a robust, All-Payer Claims Database?

Answer. The all-payer claims database could be used to test AI products on diverse populations of patients. For example, AI product performance can be closely examined across demographic subgroups, disability status, socioeconomic status, and more. Given that the database is comprehensive for all individuals who live in Rhode Island, the database could be used to ensure that AI products perform well on historically marginalized subgroups.

At the close of the Senate Finance Committee hearing, Chairman Wyden emphasized the importance of assessing bias in algorithms among Medicaid patients. Chairman Wyden noted that Medicaid barely came up during the hearing. There is a unique opportunity to leverage the Rhode Island claims database to conduct algorithmic bias analyses across both Medicare and Medicaid populations to characterize the differences in the types of biases across populations. Our team at the Duke Institute for Health Innovation would be happy to discuss ways of conducting health equity impact assessments for AI and emerging technologies.

AI can be used to normalize and harmonize data across different payer groups within the all-payer claims database. Assuming that there are inconsistencies in how different health-care payers structure and transmit data, novel large language models (LLMs) can be used to map inconsistent data representations to standard

terminologies and ontologies. This type of AI-supported data normalization could rapidly enhance the speed and efficiency with which public health and population health studies could be conducted.

The all-payer claims database can be used to develop and validate AI products that model long-term, distal outcomes across the life course. For example, patients who have consistently lived within Rhode Island for decades could have continuous claims data captured in the database. This is a unique asset that could be used to model medical conditions that progress over decades, such as genetic conditions, mental illness, toxin exposures, traumatic brain injuries, HIV, and chronic disease.

Question. Are you working in the development of AI with some of the medical specialty organizations (orthopedics, cardiologists, et cetera)? Are specialty organizations a useful place for benchmarking, approval, or accrediting best AI practices? And if they are, do you have any good examples of a medical specialty association that is being particularly forward and helpful at looking for the best uses of AI within the specialty?

Answer. The American Medical Association is an ecosystem partner to Health AI Partnership and has been a critical partner helping inform and develop AI best practices. The AMA published initial policies related to augmented intelligence in 2018: <https://www.ama-assn.org/system/files/2019-08/ai-2018-board-report.pdf>. More recently, the AMA published principles for AI development, deployment, and use: <https://www.ama-assn.org/system/files/ama-ai-principles.pdf>.

Some other examples of specialty society activities include:

- The American Academy of Family Physicians has done phenomenal work examining the role of AI for documentation and clinical review and summarization: <https://www.aafp.org/family-physician/practice-and-career/managing-your-practice/health-it/innovation-lab.html>.
- The American College of Radiology runs the AI Central Database that aggregates information from hundreds of FDA-cleared medical devices to assist with product procurement decisions: <https://aicentral.acrdsi.org/>.
- Interdisciplinary teams of experts are required to evaluate AI products for safety, efficacy, and equity impacts. AI applications in health care are socio-technical systems and alter the way patients and clinicians interact. In addition to clinical and technical experts, social scientists, lawyers, ethicists, and patients should be involved in AI product evaluations and evaluations should be conducted within the context of use. Many AI applications also span across clinical domains and care delivery settings. For example, while specialty groups often feel the pain point of poorly treated medical conditions, AI product integration often requires intervention upstream from clinical interactions with specialists. The AI product often needs to be integrated in primary care settings or be patient-facing to prompt preventive action. We write about the need to align clinicians across the care delivery spectrum as a primary opportunity to improve AI product adoption on the front lines: <https://sloanreview.mit.edu/article/ai-on-the-front-lines/>.

Question. Rhode Island has two, unusually good Accountable Care Organizations (ACOs).

In what ways could artificial intelligence be used to support the ACO program?

Answer. ACOs are uniquely positioned to benefit from AI product use, because financial incentives are aligned to prevent medical complications. AI products can predict events before they happen and give health care delivery organizations (HDOs) an opportunity to intervene early. Delivery of the early intervention may require costs associated with running the AI product, personnel effort to review high-risk patients, and up-front payment for medical interventions. But these payments can be more than recuperated if the AI intervention prevents a costly complication. For this reason, ACOs are well poised to prioritize use cases for AI, test AI products in real-world settings, and diffuse AI products that create real-world value.

ACOs can also use AI to lower costs by considering the following:

- Leverage AI to tailor disease monitoring and surveillance to individual patient risk to move beyond coarse screening protocols. This could reduce the amount of unnecessary and low-value testing.

- Leverage AI to monitor patients at home through remote sensors. Enhanced remote monitoring can detect physiologic deterioration early to prompt telemedicine interventions and prevent progression that would result in acute care visits and hospital admissions.
- Leverage AI to route patients to appropriate care setting based on data gathered through remote sensors as well as through LLM interactions. This can help ensure that patients who need acute care have their care rapidly escalated and can also help prevent unnecessary emergency department visits and hospital admissions.
- Leverage AI to optimize medical prescriptions to lowest cost therapeutic options that are effective for patients.
- Leverage AI to assess risk of surgical complications to optimize preoperative conditioning and to route patient to most cost-effective operative setting. This can help route patients at low-risk of complication to ambulatory surgery centers and route patients at high-risk of complication to inpatient surgery units.
- Patients would need to consent to use of data for these purposes and for interactions with LLMs.

PREPARED STATEMENT OF PETER SHEN, HEAD OF DIGITAL AND
AUTOMATION FOR NORTH AMERICA, SIEMENS HEALTHINEERS

SUMMARY

- Siemens Healthineers is a leading medical technology company with more than 120 years of history and experience bringing breakthrough innovations to market that enable health-care professionals to deliver the best care for patients—from prevention and early detection, to diagnosis, treatment planning and delivery, and follow-up care. Our core portfolio includes imaging, diagnostics, comprehensive cancer care and minimally invasive therapies, augmented by AI.
- Siemens Healthineers has been working on applying AI into medical technology for more than 20 years.
- To ensure we develop reliable algorithms that are reflective of the patient populations they will be applied towards, we continually maintain a holistic view of the patient with high-quality training data. This training data is based on a balanced cohort of people of different ages, genders, ethnicities, healthy people, and those who are sick.
- We recently partnered with the American College of Radiology (ACR) to improve transparency and patient care through the launch of the Transparent-AI program. We disclose detailed product information, including training data demographics and machine specifications, to help radiologists choose tools that meet their specific patient population needs.
- AI in health care can take two dominant forms—AI for operational or workflow improvements that help reduce physician burden and improve patient experience, and AI for clinical services. We refer to clinical AI as Algorithm Based Healthcare Services (ABHS), which are analytical services delivered by FDA-cleared devices that use AI, machine learning or other similarly designed software to produce clinical outputs for physicians to use in the diagnosis or treatment of disease.
- Siemens Healthineers has over 80 FDA-cleared products on the market that represent groundbreaking innovations for patients. One of our cleared products, AI-Rad Companion¹ is our dominant AI platform that highlights, characterizes, measures, and reports clinical abnormalities to aid the clinician in formulating a diagnosis and treatment.
- AI has enormous potential to improve access to care, diagnose disease faster and more precisely, and enable physicians to make treatment decisions based on comprehensive access to patient data in real-time.

¹General Availability Disclaimer for AI-Rad Companion: AI-Rad Companion consists of several products that are (medical) devices in their own right, and products under development. AI-Rad Companion is not commercially available in all countries. Its future availability cannot be ensured.

Chairman Wyden, Ranking Member Crapo, and distinguished members of the committee, on behalf of Siemens Healthineers, our 17,000 employees in the U.S., and approximately 71,000 employees in over 70 countries globally, thank you for the opportunity to testify today on the topic of artificial intelligence (AI) in health care. My name is Peter Shen, and I am the North America head of digital and automation for Siemens Medical Solutions USA, Inc.—also known as Siemens Healthineers. My career focus is on the introduction of new and emerging technologies in the health-care market, including in imaging, data analytics, digital ecosystems, and AI. My degree is in biomedical engineering and mathematical sciences from Johns Hopkins University.

Siemens Healthineers is a leading medical technology company with more than 120 years of history and experience bringing breakthrough innovations to market that enable health-care professionals to deliver the best care for patients—from prevention and early detection, to diagnosis, treatment planning and delivery, and follow-up care. Our core portfolio includes imaging, diagnostics, comprehensive cancer care and minimally invasive therapies, augmented by AI. We focus on addressing the deadliest diseases impacting the United States (U.S.), including cancer, neurovascular, neurodegenerative, and cardiovascular diseases. We partner with more than 90 percent of providers in health care and in addition to the medical devices we provide, we also work to address population growth and chronic disease prevalence, health-care workforce shortages and lack of access to care in underserved areas throughout the U.S., and globally. Given the depth and diversity of our product portfolio, we have the distinction of being the only medical technology company in the world capable of end-to-end cancer care—from diagnosis and screening to treatment and survivorship. This is a responsibility we take very seriously, and we keep patients at the center of everything we do.

We are committed to creating jobs in the U.S. and fostering community engagement. Our U.S. headquarters is in Malvern, PA. Our global headquarters for diagnostics is in Tarrytown, NY, and we have laboratory diagnostics manufacturing facilities that serve customers worldwide in both Walpole, MA and Glasgow, DE. Our global headquarters for molecular imaging is in Hoffman Estates, IL. Cary, NC is home to our training center, where we train thousands of engineers annually, including active service members. Our AI research and development team is housed in Princeton, NJ. Our Varian business is headquartered in Palo Alto, CA. We also have manufacturing, engineering, and research and development sites in Washington, Indiana, Tennessee, Nevada, and Colorado.

Each day, an estimated 5 million patients benefit from our 600,000 cutting-edge technologies and services worldwide. Data, digitalization, and AI to improve patient care is at the core of the work we do every day, and who we are as a company.

SIEMENS HEALTHINEERS AI EXPERIENCE AND ALGORITHM DEVELOPMENT

Siemens Healthineers has been working on applying AI into medical technology for more than 20 years. At our Big Data Office in the U.S., we created and maintain one of the most powerful supercomputing infrastructures dedicated to developing algorithms. This infrastructure allows our research scientists to collect, prepare and organize correct and secure medical data—including more than 2.1 billion curated images from more than 200 clinical providers and partners—needed to train and deliver accurate AI.

From its inception, we created and maintain a quality assurance process, which involves clinical validation to both understand the treatment outcomes associated with the curated data as well as guarantee the data being used to train our algorithms is accurate for diagnosing and treating disease. To ensure we develop reliable algorithms that are reflective of the patient populations they will be applied towards, we continually maintain a holistic view of the patient with high-quality training data. This training data is based on a balanced cohort of people of different ages, genders, ethnicities, healthy people, and those who are sick. From the inception of data collection, we work to build algorithms that are reliable, accurate, unbiased, and protect the patient.

We take great pride in the work we do to develop reliable AI and have company-wide guard rails for AI that I have included in an addendum to this testimony. In addition, we have recently partnered with the American College of Radiology (ACR) to improve transparency and patient care through the launch of the Transparent-AI program. We disclose detailed product information, including training data demographics and machine specifications, to help radiologists choose tools that meet their

specific patient population needs. ACR's public website includes comprehensive information on our FDA-cleared AI imaging products.

Partnering with physicians is essential to the adoption of AI, and its ability to be a powerful clinical tool to drive better patient outcomes.

REGULATION

Our algorithms go through a regulatory approval process with the Food and Drug Administration (FDA). We follow all AI/Machine Learning (ML)-enabled medical device regulatory requirements for premarket review and postmarket surveillance to ensure the safety and efficacy of our devices. We also engage with the FDA regularly on AI/ML and provide feedback on ways to ensure the continued safe and effective application of these technologies. In this regard, our AI is distinct from unregulated AI products.

With the rapid acceleration in development and innovation of AI, the need for the regulatory environment to be able to balance safety, effectiveness, as well as update and improve functionality, without hampering innovation and adoption is critical. While we believe the current regulatory framework is sufficient to support AI innovation, we support the continuation of flexibility in the approval process, as a one-size-fits-all approach could seriously inhibit the potential of AI, as well as efforts to facilitate global harmonization and the development of appropriate international consensus standards.

Additionally, Siemens Healthineers recognizes the importance of continuing to address unintentional potential bias in AI. We feel that these concerns are currently addressed for applications in medical devices and mitigated under existing risk management processes, quality systems, and compliance with regulatory requirements from the FDA and other regulators.

THE PATIENT JOURNEY

The patient journey is at the heart of Siemens Healthineers AI work. AI is already improving care for patients. Patients undergoing a CT (Computed Tomography) scan for lung cancer screening can be better positioned in the CT scanner to help optimize the resulting generated images, while minimizing the time the patient spends in the scanner. This is done by AI that is built into our CT scanner technology that allows our machines to identify human anatomy. Radiologists reviewing the resulting lung cancer screening CT images can utilize our AI-guided computer software as a companion to the clinician to identify small nodules and other abnormalities, including the ability to measure the density and characterize the size of suspicious nodules that were previously not possible to visualize without the assistance of AI.

Suspicious lung nodules diagnosed to be cancerous by the clinician can potentially be treated by radiation therapy. To minimize the risk that healthy tissue around the cancer is not unnecessarily radiated, radiation physicists create a radiation treatment plan, which includes the tedious task of manually drawing the unique contours of the cancerous tumor. This manual contouring potentially delays the time to treatment for the patient. Our AI-enabled auto-contouring software can automatically detect these contours of the cancerous area, significantly speeding up the patient's time to treatment and potentially eliminating extraneous treatments.

Utilizing AI at each point in the process to screen, diagnose and treat lung cancer can reduce the time to treatment. This allows for a reduction in patient stress and anxiety, more precise and faster diagnosis, and more specialized treatment that we believe will improve patient outcomes.

ALGORITHM-BASED HEALTH-CARE SERVICES (ABHS)

AI in health care can take two dominant forms—AI for operational or workflow improvements that help reduce physician burden and improve patient experience, and AI for clinical services. We refer to clinical AI as Algorithm-Based Health-care Services (ABHS), which are analytical services delivered by FDA-cleared devices that use AI, machine learning or other similarly designed software to produce clinical outputs for physicians to use in the diagnosis or treatment of disease. They provide quantitative and qualitative analyses, including new, additional clinical outputs that detect, analyze or interpret data to improve screening, detection, diagnosis and treatment. ABHS are developing rapidly and represent an additional service provided to the patient to deliver the best care possible. These are clinical uses of AI that have a separate and distinct place within the health-care AI conversation.

Siemens Healthineers has over 80 FDA-cleared products on the market that represent groundbreaking innovations for patients. One of our cleared products, AI-Rad Companion² is our dominant AI platform that highlights, characterizes, measures, and reports clinical abnormalities to aid the clinician in formulating a diagnosis and treatment. This ABHS supports physician decisions in diagnosing disease based on imaging scans.

For instance, ABHS can be an important service when used to diagnose neurodegenerative diseases. Brain morphometry, or changes in brain volume over time, can be a powerful predictor of neurodegenerative diseases, such as Alzheimer's, but neurologists are challenged with needing actionable, patient-specific volumetric data to diagnose and treat the patient more accurately. AI-Rad Companion (AIRC) Brain MR can automatically segment different structures of the brain on an MRI image, measure their volumes, and compare these volumes to data in a brain database from the Alzheimer's Disease Neuroimaging Initiative (ADNI). The AI-Rad Companion Brain MR feeds these comparative results into a report where deviations in volume from the norm are highlighted, enabling the radiologist to provide additional quantitative information to the neurology department, so that they can make a more accurate and informed diagnosis and treatment—resulting in better patient outcomes. AIRC Brain MR provides an additional service to the patient, and the physician receives essential information to diagnose neurodegenerative diseases that would be impossible without AI.

Another example of the benefit of ABHS is particularly relevant when discussing prostate cancer. Traditionally, a urologist identifies suspected areas of prostate cancer by manually reviewing written reports and pictograms of the prostate provided by radiology and as needed, acquires tissue samples from the areas in question using ultrasound-guided biopsy. We are developing an algorithm which is planned to be part of the AI-Rad Companion product family, which will automatically segment suspect areas of the prostate and characterize and measure suspicious lesions in the prostate from MRI images. This qualitative and quantitative analysis may support the urologist's decision on whether a tissue biopsy is additionally required for diagnosis or if such invasive procedure can be avoided, which is significant in managing a prostate cancer patient's well-being and minimizing unnecessary costs within the health system. This ABHS takes much of the gray area involved with prostate cancer, particularly when it comes to active patient monitoring, and provides a health-care service through data that the physician would not otherwise have to allow a more informed diagnosis and treatment decision.

These Siemens Healthineers AI healthcare services provide clinicians with otherwise unavailable quantitative and qualitative clinical data that allows them to make a more informed decision, resulting in better patient outcomes. However, these ABHS are not consistently reimbursed by the Centers for Medicare and Medicaid Services (CMS). Guaranteeing a consistent reimbursement process would empower providers to invest in AI confidently, ensuring their services are appropriately reimbursed. Without this financial support, these providers will face difficulties in embracing and integrating AI technologies, ultimately potentially denying revolutionary services to patients.

CURRENT AI ADOPTION CHALLENGE AND SOLUTION

CMS stated in its CY 2023 OPPTS/ASC final rule that, "Novel and evolving technologies are introducing advances in treatment options that have the potential to increase access to care for Medicare beneficiaries, improve outcomes, and reduce overall costs to the program." Furthermore, CMS also asked for input and suggestions from the public on a specific payment approach they might use for these services as AI becomes more widespread across health care.

While CMS has recognized the value and the complex nature of ABHS, the agency's reimbursement decisions have not uniformly and consistently ensured appropriate levels of payment for these services. This inconsistent, unpredictable approach stifles adoption by providers, especially in rural and underserved areas, and therefore, restricts patient access to new and innovative diagnostic tests and treatments. We support a solution that ensures a predictable and consistent approach by CMS—an approach that recognizes the costs to develop and integrate AI into the

²General Availability Disclaimer for AI-Rad Companion: AI-Rad Companion consists of several products that are (medical) devices in their own right, and products under development. AI-Rad Companion is not commercially available in all countries. Its future availability cannot be ensured.

clinical setting and reimburses for the distinct service that provides otherwise unavailable quantitative and qualitative clinical data, with a temporary and separate payment, for 5 years, based on manufacturer-supplied cost data. We believe this will allow for better adoption of AI to help CMS collect more data to evaluate the overall value of ABHS to patients. We also believe that more data will demonstrate ABHS's ability to increase access to care and improve outcomes for all patients.

THE FUTURE OF AI IN HEALTH CARE

AI has enormous potential to improve access to care, diagnose disease faster and more precisely, and enable physicians to make treatment decisions based on comprehensive access to patient data in real time. Siemens Healthineers is researching a patient companion tool to synthesize this data and apply AI to look for patterns and detect the potential for disease much earlier. In addition, we are working to create a digital twin of the patient that would allow a physician to perform an interventional procedure, say for a heart procedure, on a digital replica of a patient's heart to test how that patient will react and respond to a specific course of treatment before it is applied to the individual. The digital twin will minimize unintended consequences and provide more personalized, precision medicine for the patient. We are excited about what the future holds for AI in health care.

CONCLUSION

In closing, thank you for the opportunity to testify before you today on AI in health care. While there are many forms of AI applications in health care to reduce physician burnout and streamline operational complexities, we believe the highest value of AI in health care comes in the form of ABHS, and that this will revolutionize health-care services for patients. Siemens Healthineers is a market leader in researching and training AI in medical technologies and welcomes the opportunity to continue this discussion. It is critical that we all work together to ensure we create trust with consumers and build ethical, transparent, and accessible AI in health care to improve patient outcomes. Again, thank you for the opportunity to testify. I look forward to your questions.

ADDENDUM

We use a set of guard rails to guide the way we develop and implement AI in health care:

- We believe that health-care professionals, backed up by AI solutions, make a strong team.
 - Our AI solutions learn from the best: Siemens Healthineers collaborates with a huge network of world-class clinicians, where we combine our research and development (R&D) capabilities with our customers' clinical expertise. The results of this collaborative process are powerful, clinically proven AI companions for decision-making that help to provide better patient care at lower cost. Humans and artificial intelligence have vastly different abilities. We believe that the future of medicine lies in combining the strengths of these capabilities. Such systems will provide health-care professionals with tools to meet the rising demand for diagnostic imaging and actively shape the transformation of radiology into a data-driven research discipline. Moreover, AI algorithms are expected to help speed up clinical workflows, prevent diagnostic errors and reduce missed billing opportunities, thus enabling sustained productivity increases.
- We believe the level of autonomy of AI solutions needs to be balanced with ethical expectations and human values.
 - Societies are currently discussing the extent to which AI solutions could be a vital part of everyday human life. Depending on the area of life, society allows and strives for lower or higher levels of autonomy. In this regard, health care is a special area, as patients benefit from and rely on the trusted doctor-patient relationship. A high degree of autonomy of an AI solution substantially impacts this relationship. In health-care areas, where the personal and trusted patient-doctor relationship is key to the success or course of the treatment, we believe that the autonomy of AI solutions needs to be well-balanced. Therefore, we develop AI solutions only for areas where they are ethically acceptable and beneficial to humankind and society.

- We develop AI solutions to support patients' desires for more personalized medicine.
 - An increasing choice of personalized therapies is leading to significantly improved outcomes in oncology, but personalized medicine is also gaining traction in other application areas. For physicians, however, it is becoming more and more challenging to keep abreast of the constantly expanding treatment options. With our AI solutions, we enable physicians to make more accurate diagnosis and treatment choices, based on comprehensive patient data and the ever-advancing wealth of medical knowledge. With our vision of the "Health Digital Twin" as a constantly updated virtual model of the human body, we strive to develop the next generation of systems for personalized medicine.
- We believe data handling in health-care needs to focus on the individual.
 - We support patients, so they can share their health data safely and securely with physicians in health systems. Our e-health solution creates a decentralized electronic health record that enables patients to make their longitudinal health data accessible to physicians. The patient is in control and decides who to share their data with. We promote the vision of a "Health Digital Twin" in health care, which models and represents a human body based on a multitude of datasets like body composition and vital parameters. For both patients and healthy people, their digital twin will help physicians to diagnose complex systemic diseases earlier and find the best treatment available for the patient's given condition.
- We strive to develop AI solutions for both healthy people and sick people.
 - Our current portfolio focuses on diagnosing and treating patients. Yet, we believe that stewardship for a patient starts with prevention, and the predictive power of AI offers a wealth of opportunities for us to help people stay healthy. In the future, we want to extend our portfolio to support health systems in their transformation from caring for the sick to proactively caring for the well.
- We work passionately to make AI solutions accessible to patients everywhere.
 - At Siemens Healthineers, we believe that every human being has the right to access high-quality health care, regardless of location, age, and social circumstances (in line with Art. 27 (1) Declaration of Human Rights "right to progress"). Thus, we support the United Nations' 3rd Sustainable Development Goal (SDG), which ensures healthy lives and promotes well-being for all at all ages. By providing powerful AI solutions, we contribute to better and more personalized health care that is accessible around the globe.
- We believe AI development needs to be transparent.
 - We openly communicate insights into underlying technology, training/test datasets, and quality assurance for our AI solutions. We carefully compile training and test datasets which we document to allow traceability and transparency. Specifically, we strive to free our data from bias and prejudice to enable equal treatment for all people.
- We measure ourselves against the highest scientific standards.
 - We aim to improve clinical outcomes with state-of-the-art technologies. We do not fuel technological hype; instead, we invest in science to improve technology and establish new standards. Our world-class scientists therefore critically evaluate and thoroughly assess our AI solutions with carefully designed evaluation studies for the respective target populations.
- We speak honestly about the capabilities of our AI solutions.
 - We are aware of the capabilities and limitations of our AI solutions and share these insights with our customers and users in order to promote the setting of realistic expectations. Expectations of any technical system need to be realistic to prevent false hopes, misunderstandings, and errors in judgment. Health-care professionals need to be aware of the capabilities of an AI solution, so that they can make an informed decision in line with applicable best practices and guidelines and advise patients accordingly.

Data Privacy—we believe that to fully realize the potential of digital transformation, people need maximum confidence in the processes, institutions, and technologies used.

At Siemens Healthineers, our data vision is, “we use data responsibly to develop innovations in health care to help people live healthier and longer lives.” This vision has given rise to a set of data principles that guide our handling of very sensitive health data and the development of today’s and tomorrow’s digital health solutions:

- We use data for the benefit of the individual.
 - The purpose of our company is to advance human health. People should benefit from data-driven medical innovations through the prevention of sickness and best-in-class procedures and treatment. We invest in data-driven health solutions because we support the patient’s desire for personalized high-precision medicine to live a healthier and longer life.
- We use data to drive health-care innovation.
 - Data will become the key enabler for innovations in digital health care. Data-driven innovations are essential for medical research and progress. Our tailored and responsible use of data enables us to fill our innovation pipeline, push data-driven medicine and develop innovative procedures for patients.
- We are trustworthy and ethical in our handling of data.
 - We only use data in a purpose-bound manner to develop medical innovations and to enable our data-driven products to perform according to their specified performance capabilities. We treat data responsibly, reliably, and securely.
- We apply proven and high data privacy standards worldwide.
 - We believe that trust and accountability are basic pillars for responsible data privacy management. Consequently, we apply high data privacy standards worldwide. Fundamental legal principles of the GDPR—including the legitimacy and lawfulness of data processing, purpose limitation, the need-to-know principle, data avoidance and data economy—are mandatory for Siemens Healthineers worldwide based on internal directives. In addition, we apply proven technical standards and organizational measures to ensure data security, authenticity, and confidentiality. Our ISO-certified cybersecurity management system follows a holistic approach and integrates information security management (ISO 27001) and privacy information management (ISO 27701).
- We support the advancements that enable individuals to have sovereignty and transparency over their data.
 - Every person should have sovereignty over their own health data. This includes transparency on what data is used on what basis and for what purposes, and the right to grant or revoke consent to the use of one’s own data. This right should also include the freedom to donate one’s personal data for the purpose of conducting research, advancing progress, and improving health-care solutions. The processing of health data in private-sector research and development work also contributes significantly to advancing medical and technical progress. To safeguard this valuable contribution, we believe that private-sector research is also subject to the privilege of research, and that the development of medical devices or artificial intelligence that facilitate(s) improvements in the early detection or treatment of illnesses, for instance, also serves the public interest and public health. We promote trust throughout society and among all patients for the application of digital technologies and support the exercising of their rights accordingly.
- We leverage data as a strategic asset.
 - Driving digitalization and promoting value creation from data are essential to advancing medical progress and providing efficient, high-quality health care. Leveraging this potential of data is strategically important to us. Besides developing data- and software-driven solutions for supporting decision-making, we continuously pursue efforts to further develop our portfolio by automating devices and workflows and expanding our use of predictive maintenance. The interoperability and connectivity

of our products and solutions accelerates this development into a platform-oriented business.

- We use state-of-the-art technology to protect data.
 - We offer a state-of-the-art portfolio of secure products, cybersecurity services and consulting that helps to ensure optimum protection. We continuously improve our systems and processes and train our teams in aspects of cybersecurity and data protection to maintain a consistently high level of threat awareness is. Our engineering practices include a secure development lifecycle (SDL) to ensure that high cybersecurity standards are implemented for every product and solution. Examples of our core development principles are the implementation of privacy by design and privacy by default.
- We support open standards for data interoperability.
 - The key to data-driven health-care innovations is the ability to interconnect various health datasets. It is only through data integration and data interoperability that the value of data can be fully utilized. We strongly support the standardization of health-care data and data sharing. When designing our solutions, we aim to systematically include standardized interfaces such as DICOM5, FHIR6, and increasingly uniform APIs7.
- We invest in trustful partnerships to access data.
 - Efforts to improve medical knowledge and to advance data-driven health-care solutions depend on having rights to access health data from diverse, genuine sources. We believe that providing fair access to relevant data by all health-care stakeholders and using this data responsibly to our mutual benefit will contribute to advancing medical progress. We therefore build our data-related partnerships on fairness and transparency.

QUESTIONS SUBMITTED FOR THE RECORD TO PETER SHEN

QUESTIONS SUBMITTED BY HON. CHUCK GRASSLEY

Question. We spent about \$4.5 trillion in health care last year. A key area of waste in our health-care system is medical errors, failure of care delivery, and over-treatment. Some estimates suggest we waste \$205 billion to \$425 billion each year due to medical errors and waste.

What potential does deploying artificial intelligence to review clinical decisions have in reducing medical errors and waste? Are you aware of existing practices and can we scale them?

Answer. Algorithm-Based Health-care Services (ABHS) produce qualitative and quantitative clinical findings that support physician diagnostic and therapeutic clinical decisions with increased sensitivity, specificity, and accuracy. ABHS helps physicians recognize potential clinical ailments earlier in diagnosis and minimizes additional unnecessary diagnostic tests, such as identifying, characterizing, and quantifying coronary calcification in a patient undergoing a routine chest CT screening exam.

Question. According to a Mercatus Institute analysis, the Department of Health and Human Services is home to over 42,000 Federal Government regulations for health care, including over 16,000 regulations at the Centers for Medicare and Medicaid Services and over 13,000 regulations at the Food and Drug Administration.

Should we add more Federal regulations for artificial intelligence, or do existing regulations protect safety and promote good governance without stifling innovation?

Answer. Siemens Healthineers algorithms go through a regulatory approval process with the Food and Drug Administration (FDA). We follow all AI/Machine Learning (ML)-enabled medical device regulatory requirements for premarket review and postmarket surveillance to ensure the safety and efficacy of our devices. We also engage with the FDA regularly on AI/ML and provide feedback on ways to ensure the continued safe and effective application of these technologies. In this regard, our AI is distinct from unregulated AI products.

We believe that with the rapid acceleration in development and innovation of AI, the need for the regulatory environment to be able to balance safety, effectiveness, as well as update and improve functionality, without hampering innovation and adoption is critical. While we believe the current regulatory framework is sufficient to support AI innovation, we support the continuation of flexibility in the approval process, as a one-size-fits-all approach could seriously inhibit the potential of AI, as well as efforts to facilitate global harmonization and the development of appropriate international consensus standards.

Additionally, Siemens Healthineers recognizes the importance of continuing to address unintentional potential bias in AI. We feel that these concerns are currently addressed for applications in medical devices and mitigated under existing risk management processes, quality systems, and compliance with regulatory requirements from the FDA and other regulators.

Question. You work for Siemens, and that's a large company. Do you think a small hospital or entrepreneur can navigate the existing Federal regulatory regime to get their artificial intelligence product to market?

Answer. As part of the MedTech industry, our Algorithm-Based Health-care Services (ABHS) go through an established, proven regulatory approval process with FDA. While we believe the current regulatory framework is sufficient to support AI innovation from all MedTech vendors, we support the continuation of flexibility in the approval process, as a one-size-fits-all approach could seriously inhibit the potential of AI in health care.

What we see as a big problem for small hospitals and rural health-care facilities is an inconsistent and unreliable reimbursement approach from CMS for FDA-cleared clinical AI products. This inconsistent reimbursement approach could have significant impact on adoption and access to care for rural health-care facilities, resulting in patients not having access to innovative services that help support more informed diagnosis and treatment decisions.

Question. At the House Energy and Commerce Committee hearing last November, you described Siemens Healthineers' AI Office of Big Data and mentioned you've created and maintained a transparent quality assurance process.

Are there policy lessons we can learn from Siemens Healthineers' efforts to ensure transparency and quality assurance in using artificial intelligence for health care?

Answer. Siemens Healthineers has been working on applying artificial intelligence in medical technology for more than 20 years. Over those 20 years we have found that keeping patients at the heart of every step of the process, and building trust through a transparent end-to-end approach, are the keys to reliable, quality AI services. We encourage others in the private sector to adopt principles to ensure transparency and quality in the training of algorithms, and we believe our principles are a good starting place for that discussion.

At our Big Data Office in the U.S., we've built one of the most powerful supercomputing infrastructures dedicated to developing AI in health care. This allows our research scientists to collect, prepare and organize correct and secure medical data—including more than 1.8 billion curated images from more than 200 clinical providers and partners—needed to train and deliver accurate AI algorithms. From its inception, we have created and maintained a quality assurance process, which involves clinical validation to both understand the treatment outcome associated with the curated data as well as guarantee the data being used to train the AI algorithms is accurate for diagnosing and treating disease.

Additionally, to ensure we develop reliable AI algorithms that are reflective of the patient populations they will be applied towards, we continually maintain a holistic view of the patient with high-quality AI training data. This training data is based on a balanced cohort of people of different ages, genders, ethnicities, healthy people, and those who are sick. We work from the inception of AI development to build algorithms that are reliable, accurate, unbiased, and protect the patient.

We are proud of the work we do to develop reliable, quality AI and have developed company-wide guard rails for AI that were included in the addendum of our written testimony.

QUESTIONS SUBMITTED BY HON. MARIA CANTWELL

Question. One of the most difficult questions we're struggling with right now is predicting the impact that AI will have on our workforce. When we have previously seen rapid technological advances, people have often worried about job loss. However, when it's all said and done, the technological advances actually created more jobs.

It's undeniable that Washington State needs more health-care workers. The Washington State Hospital Association reports that the State's hospitals need to hire over 6,000 more nurses to meet their staffing needs. Maybe AI will help with job creation, but I'm not so sure. As Mustafa Suleyman, cofounder of Google's Deep Mind, commented at Davos this year, AI is basically a labor-replacing tool. While I believe that AI has tremendous potential to address our health-care workforce shortage, I'm very concerned with the potential negative impact that AI could have on our health-care work force.

At the end of the day, computers and algorithms cannot replace the logic and reasoning capabilities of a human who has gone to medical school or received a nursing degree and gained experience through practice. Artificial intelligence also cannot replicate, for example, a nurse who employs empathy and human characteristics while comforting a terminal patient. There is a role for AI in improving the resiliency of the health-care workforce, but AI should not completely replace human presence and decision-making.

In your opinion, when is it appropriate to employ artificial intelligence to fill the gaps that the health-care workforce shortage has created? When is it also not appropriate?

Answer. Artificial intelligence can be an important tool to streamline hospital operations and efficiencies, for instance, to ensure optimal use of scheduling for operating rooms and staffing needs. We believe there is a role for AI in supporting administrative tasks to help alleviate workforce shortages. However, we believe clinicians must stay at the center of patient engagement—from screening and diagnosing to treating disease. It is not appropriate to use AI to replace physician decision-making. Specifically, the physician must understand, evaluate, and determine whether to use Algorithm-Based Health-care Services (ABHS) when interacting with the patient. This is an additional service available to the physician but must not be used in lieu of that physician. We believe that health-care professionals, backed up by AI solutions, make a strong team for the patient.

Question. Your company, Siemens Healthineers, partnered with the American College of Radiology to foster innovation and improve innovative patient care. How does that project balance the need to fill workforce shortage gaps while protecting the role of humans in health-care settings?

Answer. Our partnership with the American College of Radiology (ACR) demonstrates the importance of working closely with specialty organizations in the successful adoption of AI in health care. Through the launch of the Transparent-AI program we disclose detailed product information, including training data demographics and machine specifications, to help ACR members understand and choose efficient AI tools that meet their specific patient population needs while also addressing their own workforce challenges.

Additionally, we believe health care is a special area, as patients benefit from and rely on the trusted doctor-patient relationship. A high degree of autonomy of an AI solution substantially impacts this relationship. In health-care areas, where the personal and trusted patient-doctor relationship is key to the success or course of the treatment, we believe that the autonomy of AI solutions needs to be well-balanced. Therefore, we develop AI solutions only for areas where they are ethically acceptable and beneficial to humankind and society.

QUESTIONS SUBMITTED BY HON. JOHN CORNYN

Question. While CMS has recognized the value and the complex nature of Algorithm-Based Health-care Services (ABHS), the agency's reimbursement decisions have not uniformly and consistently ensured appropriate levels of payment for these services. This inconsistent, unpredictable approach stifles adoption by providers, especially in rural and underserved areas, and therefore, restricts patient access to new and innovative diagnostic tests and treatments. We support a solution that ensures a predictable and consistent approach by CMS—an approach that rec-

ognizes the costs to develop and integrate AI into the clinical setting and reimburses for the distinct service that provides otherwise unavailable quantitative and qualitative clinical data, with a temporary and separate payment, for 5 years, based on manufacturer-supplied cost data. We believe this will allow for better adoption of AI to help CMS collect more data to evaluate the overall value of ABHS to patients. We also believe that more data will demonstrate ABHS's ability to increase access to care and improve outcomes for all patients.

You state in your testimony that CMS coverage and reimbursement is an “inconsistent, unpredictable approach” and that you “support a solution that ensures a predictable and consistent approach” by the agency. I support a solution as well. As we have considered it, part of the reason I see for the inconsistent and unpredictable approach is that Medicare’s payment systems act independent of one another and often try to fit innovative products like AI into their coverage systems. Private insurance, however, seems much more in the business of building new coverage and reimbursement approaches for new innovative products to get more savings/value out of them.

Do you see this siloed approach to operating the Medicare program as part of the problem like I do?

Answer. The challenge lies in ensuring Medicare payment systems are inter-related. One notable instance demonstrating progress in this direction is evident in the CY 2024 Medicare Physician Fee Schedule (PFS) final rule issued by CMS. This discussion highlighted the authority and appropriateness of cross-walking payments for certain services from the Hospital Outpatient Prospective Payment System (OPPS) to PFS. It reflects a promising step towards fostering interrelatedness within Medicare’s payment mechanisms.

However, to fully capitalize on this progress, it’s imperative that CMS consistently exercises this authority when appropriate. By doing so, CMS can help mitigate the fragmentation and inconsistencies that currently characterize Medicare’s payment landscape. This approach would not only enhance the efficiency of payment processes but also ensure equitable access to innovative services, such as those involving AI, across different health-care settings.

Question. An article in *Health Affairs* by Robert Horne¹ theorized that one way to address this problem of inconsistent unpredictable approach is to reverse the CMS coverage and payment process. Instead of focusing on coverage within a specific payment system, establish program-wide coverage and reimbursement with CMS leadership and then work to imbed these new arrangements with individual payment systems. This would allow CMS as a program to consider and capture more of the value from innovative product designs like AI as well as the means of pushing consistency down to the payment systems.

What do you think about this approach?

Answer. We agree that “ineffective approaches to payment can also lead to increased program expenditures without additional benefits.” Therefore, we strongly encourage CMS to adopt a payment framework that adequately accommodates AI, particularly products falling under the definition of Algorithm-Based Health-care Services (ABHS). Pursuing a modern, comprehensive, and reliable reimbursement pathway that acknowledges the value of AI is critical; however, more data is required to determine metrics that could be universally applied across all payment systems. Presently, the emphasis is on CMS utilizing its existing pathways consistently, while broader endeavors persist in addressing AI coverage and payment on a larger scale.

QUESTION SUBMITTED BY HON. JOHN THUNE

Question. AI has the potential to improve certain aspects of administrative processes and practice in clinical settings that could make hospitals and physicians more efficient and improve safety and quality.

How can the Medicare program appropriately value these services and account for increased efficiency and reduced costs?

¹ <https://www.healthaffairs.org/content/forefront/digital-era-payment-reform-key-shaping-modern-medicare-program>.

Answer. Although AI technology undoubtedly holds promise for enhancing efficiency in health-care delivery, our current emphasis lies on AI tools capable of furnishing actionable clinical data to support physicians in decision-making processes, ultimately leading to improved patient outcomes. Medicare's ability to accurately assess and value these services is crucial. This involves recognizing not only the costs associated with developing and integrating AI into clinical settings but also acknowledging the unique value proposition of AI-enabled services, which provide access to quantitative and qualitative clinical data that would otherwise be unavailable. By appropriately reimbursing for these distinct services, Medicare can incentivize the adoption of AI technology and facilitate its integration into routine clinical practice, thereby driving advancements in patient care and health-care delivery.

QUESTIONS SUBMITTED BY HON. SHELDON WHITEHOUSE

Question. Rhode Island has several health-care enterprises that we are very proud of, including the State-wide Rhode Island All-Payer Claims Database (RI APCD).

What role can AI play with regard to maximizing the utility of a robust, All-Payer Claims Database?

Answer. Generative AI, which provides the potential to sort through information rapidly to identify certain characteristics and outliers, could be used with an All-Payer Claims Database to look for patterns and signal potential abnormalities that would encourage a physician or clinician follow-up. This would require strong privacy protections, a clear and distinct chain of responsibility for data ownership and protection, and cybersecurity protections. While there is great potential here, there is also great risk and consumer trust must be established prior to the exploration of AI within an All-Payer Claims Database, depending on what information AI has access to.

Question. Are you working in the development of AI with some of the medical specialty organizations (orthopedics, cardiologists, et cetera)? Are specialty organizations a useful place for benchmarking, approval, or accrediting best AI practices? And if they are, do you have any good examples of a medical specialty association that is being particularly forward and helpful at looking for the best uses of AI within the specialty?

Answer. Partnering with physicians is essential to the adoption of AI, and its ability to be a powerful clinical tool to drive better patient outcomes. As example, our partnership with the American College of Radiology (ACR) demonstrates the importance of working closely with specialty organizations in the successful adoption of AI in health care. Through the launch of the Transparent-AI program we disclose detailed product information, including training data demographics and machine specifications, to help ACR members understand and choose efficient AI tools that meet their specific patient population needs while also addressing their own workforce challenges.

Question. Rhode Island has two, unusually good Accountable Care Organizations (ACOs).

In what ways could artificial intelligence be used to support the ACO program?

Answer. Algorithm-Based Health-care Services (ABHS) produce qualitative and quantitative clinical findings that support physician diagnostic and therapeutic clinical decisions with increased sensitivity, specificity, and accuracy. In these ways, ABHS could be incredibly helpful to accountable care organization physicians by better assisting in recognizing potential clinical ailments earlier in diagnosis and minimizing additional unnecessary diagnostic tests such as identifying, characterizing, and quantifying coronary calcification in a patient undergoing a routine chest CT screening exam.

PREPARED STATEMENT OF HON. RON WYDEN,
A U.S. SENATOR FROM OREGON

This morning the Finance Committee meets to discuss the use of artificial intelligence, or AI, systems in health care, with a focus on how this technology is being used in Federal health programs like Medicare and Medicaid.

There's no doubt that some of this technology is already making our health-care system more efficient. But some of these big data systems are riddled with bias that

discriminates against patients based on race, gender, sexual orientation, and disability. It's painfully clear not enough is being done to protect patients from bias in AI.

Just as I worked to ensure tech innovation would improve patient care in the 1990s with laws empowering telemedicine and digital signatures, Congress has an obligation to encourage the good outcomes from AI and set rules of the road for new innovations to deliver better care for Americans.

Today we'll also discuss the role Congress and this committee must play in helping strike a balance between protecting innovation and protecting patients and their privacy with legislative proposals like my Algorithmic Accountability Act, which would tackle these concerns head-on.

There are a lot of reasons to be optimistic about the potential of AI to improve health care. Today the industry is facing a host of challenges, all made worse by the strain the COVID-19 pandemic put on our health-care system. There's an ongoing workforce shortage, existing providers are facing high rates of burnout, health-care costs are rising faster than wages, and there's an ever-growing gap between the care that's needed and the care actually being delivered to many Americans.

Already, AI tools are being deployed to reduce some of these pressures and ease strain on the industry and providers. Some doctors are using this technology to prepopulate clinical notes and emails to reduce workload, submit bills to insurers to reduce administrative waste, and even help with diagnostics. Primary care providers can use these tools to screen for certain diseases and connect patients with specialists for treatment, saving patients' time and money, and leading to better, more timely care.

There is no doubt these technological innovations can improve care for patients in Medicare and Medicaid, while also improving workload for providers, many of whom are already stretched thin. But addressing these challenges with new technology shouldn't mean worse patient outcomes and sacrificing patient privacy.

Unfortunately, there are clear, glaring examples of AI tools being developed with data that perpetuate racial biases, and deployed in ways that bypass important doctor expertise, leading to inadequate care for patients.

The committee is lucky to have here today Dr. Ziad Obermeyer, who in 2019 discovered harmful racial bias in an AI tool developed by the health-care company Optum—a subsidiary of UnitedHealth Group—and used by providers across the country to offer care management services.

He found that the tool, on average, required Black patients to present with worse symptoms than White patients in order to qualify for the same level of care.

This algorithm was available to thousands of doctors across the country, potentially impacting millions of patients. How could such a flawed system make its way into general use? The answer is simple: there was nobody watching. No guard rails were in place to protect patients from flawed algorithms and AI systems.

To make matters worse, the technology that insurance companies or health systems use can play a role in what care patients receive, and what services are approved or denied. And the Department of Health and Human Services does not—yet—oversee the use of these systems.

I think most of us here would agree there are many ways this technology can be used to improve health care and patient outcomes. However, as we increasingly rely on technology like AI to make decisions in every facet of our day-to-day lives, this committee has a responsibility to ensure there are guard rails in place to protect patients, particularly in Medicare and Medicaid, and I do not believe that current laws go far enough to achieve that goal.

My Algorithmic Accountability Act lays the groundwork to root out algorithmic bias from these systems. As applied to health care, my bill would require health-care systems to regularly assess whether the AI tools they develop or select are being used as intended and aren't perpetuating harmful bias.

I'll close with this: I believe the same protections in my Algorithmic Accountability Act must apply to patients in Medicare and Medicaid. Here's what's needed most: transparency in how these tools are developed and used to foster trust, and accountability for how these tools are used in health care. These tools should also preserve the privacy of patients. Lastly, these tools should further equity in health

care, not perpetuate harmful bias or disadvantage hospitals and providers who serve low-income patients or communities of color.

The Food and Drug Administration and the Office of the National Coordinator for Health IT have proposed new rules to address some of these issues. That's a step forward. But they don't go far enough. It's clear more is needed to protect patients from flawed systems that can and will directly affect the health care they receive. I look forward to working with my colleagues on the committee to identify ways we can protect patients and improve care going forward.

COMMUNICATIONS

AARP

<https://www.aarp.org/>

AARP, which advocates for the more than 100 million Americans age 50 and older, appreciates the Senate Committee on Finance’s effort to better understand the growing impact of artificial intelligence (AI) and other algorithmic tools in the delivery of health care. As major consumers of health care, older Americans will acutely experience the benefits and potential harms of AI’s expansion into all aspects of our health care system.

The expanded use of AI holds great promise for improving patient care. As noted in a recent U.S. Government Accountability Office report¹ to Congress, clinical AI tools show encouraging results in predicting health trajectories of patients, recommending treatments, guiding surgical care, monitoring patients, and supporting population health management. These exciting tools have huge potential to increase quality and efficiency in health care, reduce complexity and inefficiency in consumer interactions, and make other improvements in ways that we cannot yet comprehend.

At the same time, the growth of AI presents significant risks that could negatively impact patients in numerous ways. For example, AI could be used to augment profit-seeking behavior in the health care system by automatically denying certain claims and leaving patients with no understanding of how those decisions were made and little recourse to correct them. AI also relies on existing datasets to develop algorithmic decision tools, which can reinforce existing biases and disparities in the health care system. For instance, under-representation of minorities because of racial biases in dataset development might lead to subpar prediction results² for members of those groups. Without proper safeguards, these tools can make decisions that perpetuate those biases, reflecting historic prejudices, including against older adults.

While it is difficult to predict all of the ways in which AI will impact the health care system in the years to come, we encourage Congress to keep the experience of consumers and their safety as the highest priorities and carefully consider the following principles when developing policy.

General Principles

Fairness, transparency, and accountability should guide all uses of AI or other algorithmic tools that make consequential decisions regarding a patient’s health, coverage, or well-being. Patients must feel secure that the algorithmic tools being used to make potentially life and death decisions are:

- Fair, reliable, and accurate and not result in unjustifiable disparate impacts on people’s civil rights;
- Transparent when a consumer, or provider involved in a consumer’s care, interact with such a tool and provide clear explanations about the results; and
- Accountable, such that patients who have been adversely affected by a decision informed by an algorithmic decision tool have access to a fair and meaningful process to challenge those decisions and their outcomes.

The degree of regulation should be commensurate with the potential risk of harm to individuals and should focus on outcomes and performance standards, not the technology used.

¹ <https://www.gao.gov/products/gao-21-7sp>.

² <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC9908503/>.

Anti-Discrimination

Digital discrimination and algorithmic bias pose new challenges in health care. For example, AI may determine whether a person qualifies for a particular health care service or government benefits. The use of algorithms in these situations has the potential to amplify systemic discrimination and may be outside the scope of existing civil rights laws. Congress must ensure that consumer protection and anti-discrimination laws that apply in the context of human decisions are adapted to effectively apply to AI algorithmic decision tools.

AI tools used in the health care context should be evaluated for accuracy, reliability, and fairness prior to deployment and routinely thereafter. If it is not reasonably possible to calculate the impact of each input or factor in an algorithm on a decision, the algorithm as a whole must be evaluated to determine if it results in an unjustifiable disparate impact on a protected class. When governments utilize algorithmic tools to inform consequential decisions, these tools should be evaluated by a qualified third party for reliability, accuracy, and possible biases against people's civil rights protections. The results of these evaluations should be made public.

Similarly, strong protections should be enacted to protect against AI-based improper denials of coverage in insurance decisions and incorrect diagnosis in medical decisions. Consumers must also have peace of mind that if their insurance carrier makes an erroneous denial of coverage or their medical provider makes an incorrect diagnosis there is a clear path to recourse through an expedient and nonbiased appeals process.

Transparency and Privacy

For consumer protections to be meaningful, there must be transparency and accountability. Transparency must include providing clear, readily accessible notice to people when they are interacting with an algorithm, the intent of the algorithm, clear explanations about the results, and any relevant implications. In addition, transparency includes providing people with access to easy-to-understand explanations of the factors that contributed to a decision and the logic behind it.

At their most fundamental level, all applications of AI rely on the analysis of reams of data to detect patterns and make inferences and predictions. The challenge is establishing the guardrails that allow for data uses that bring lasting consumer benefits while still providing robust consumer privacy protections. Privacy protections should be embedded into all products and services, keep pace with changing technology and privacy standards, minimize data collection to what is needed for the product or service, and be developed with strong input from consumer stakeholders. Organizations (including private companies, nonprofits, and government entities) should clearly communicate to the consumer:

- What identified and deidentified personal information is collected, inferred, or deduced, and how it can be used, maintained, shared, or sold to others;
- Whether AI systems are applied to personal information;
- How data created from the use of AI are used, maintained, shared, or sold to others; and
- How to exercise opt-out rights.

Privacy policies regarding a consumer's data should be written in plain language and disclosed before a consumer uses a product or service. They should be clear, short, and standardized. Organizations should be required to evaluate and mitigate the privacy risks to consumers, and privacy laws and regulations should include strong enforcement mechanisms to ensure compliance. These mechanisms include strong enforcement authority, appropriate fines and penalties, and swift compliance deadlines.

Managing health-related data, in particular, presents challenges and risks that need a comprehensive framework of privacy and security safeguards to help ensure public trust. Organizations engaging with health data should be required to provide consumers with meaningful transparency, choice, and control related to their health and health-related data, and be required to:

- Ensure accuracy of health and health-related data;
- Provide consumers with the opportunity to review the health information and data held about them;
- Allow consumers to dispute and resolve the accuracy or completeness of health and health-related data;
- Provide consumers with an easy and swift process to delete data when they want to end the relationship with an organization; and

- Obtain consumers' explicit opt-in consent before selling, sharing, or trading personally identifiable health or health-related data.

Conclusion

Thank you for the opportunity to provide AARP's perspective on AI and health care. We look forward to working with you to address this important issue to ensure that all Americans receive timely, safe, and quality care as these new technologies are widely deployed.

ADVANCED MEDICAL TECHNOLOGY ASSOCIATION (ADVAMED) IMAGING

Algorithm-Based Healthcare Services

Advanced Medical Technology Association (AdvaMed) Imaging appreciates the opportunity to provide a statement for the record in response to the U.S. Senate Committee on Finance hearing entitled Artificial Intelligence and Health Care: Promise and Pitfalls. AdvaMed Medical Imaging Division represents the manufacturers of medical imaging equipment, including, magnetic resonance imaging (MRI), medical X-Ray equipment, computed tomography (CT) scanners, ultrasound, nuclear imaging, radiopharmaceuticals, and imaging information systems. Our members have introduced innovative medical imaging technologies to the market, and they play an essential role in our nation's health care infrastructure and the care pathways of screening, staging, evaluating, managing, and effectively treating patients with cancer, heart disease, neurological degeneration, COVID-19, and numerous other medical conditions. We focus our statement on how Congress can improve innovation, adoption, and access to algorithm-based healthcare services (ABHS).

ABHS are services enabled by medical devices cleared by the Food and Drug Administration (FDA) that rely on artificial intelligence (AI) or machine learning (ML), or other similar software to produce quantitative and/or qualitative outputs that clinicians can use to aid in the diagnosis or treatment of a patient's condition. ABHS provide clinical outputs that cannot be otherwise obtained by a healthcare provider and has played a growing role in improving care delivery and informing care pathways throughout the healthcare industry. For example, ABHS can identify diabetic retinopathy, and measure the caliber of coronary arteries via mechanisms and clinical outputs not otherwise available to inform patient care.

We note that ABHS occupy a unique niche in the broader discussion of AI/ML in health care, as ABHS undergoes rigorous FDA regulatory approval pathways and review by the Centers for Medicare & Medicaid Services (CMS) before such services will be covered under the Medicare program. Additionally, in contrast to other AI/ML software, ABHS are not generative AI models or used outside the supervision and guidance of a trained healthcare professional. Nor is ABHS identical to other uses of AI/ML in health care, such as software that passively monitors vital signs and improves provider workflows (*e.g.*, AI that improves clinical documentation in electronic medical records).

While all types of ABHS described by the American Medical Association's (AMA) Current Procedural Terminology Panel continues to demonstrate its immense value, less than 10 ABHS have (1) received FDA approval or clearance, (2) received a Current Procedural Terminology CPT code from the American Medical Association (AMA), and (3) received CMS coverage and payment through a unformalized case-by-case basis. The ABHS that receives Medicare coverage includes valuable diagnostic tools such as characterizing potentially cancerous lung nodules and assessing signs of liver disease.

However, only a few ABHS have received payment assignments through Medicare, which is in sharp contrast to the nearly 600 FDA approved or cleared AI/ML medical devices, many of which are ABHS. CMS is ill-equipped to handle the unique aspects of ABHS, and in fact, is looking for guidance from stakeholders on the best approach. For example, today, an FDA approved ABHS that providers can use to identify potentially cancerous colorectal lesions is not separately payable by CMS, despite the similarity it has with already covered ABHS that can characterize lung nodules. This case-by-case reimbursement approach will stifle innovation and patient access to cutting-edge diagnostic tools and treatments.

To ensure that patient access and innovation continues for ABHS, CMS must modernize its policies to account for the uniqueness of ABHS. Under the current system, access to these important services will be stunted and innovation will decline.

AdvaMed Imaging recommends that Congress encourage CMS to formalize its existing Software-as-a-Service add on payment policy, revise its New Technology Ambulatory Payment Classification (APC) policies to account for the unique aspects of ABHS, and provide ABHS with at least 5 years of consistent payments while assigned into a New Technology APC. Currently, CMS has the authority to make these policy changes.

AdvaMed Imaging further suggests that Congress define ABHS in the Social Security Act in order to better identify and distinguish ABHS from other types of AI/ML. Congress should also include ABHS as a covered and reimbursed hospital service. By doing so, Congress will promote innovation and beneficiary access to an important healthcare service that is expected to lead to earlier and more accurate diagnoses, faster treatment and improved patient outcomes.

Finally, Congress should be judicious in imposing new and unnecessary requirements for ABHS, given the demanding regulatory frameworks already in place. FDA currently applies strict requirements to the development, testing, and approval of ABHS, including a rigorous pre-market review, processes that assess ABHS performance, reliability, and safety, and ongoing monitoring and surveillance after the ABHS has been approved or cleared. ABHS developers must also adhere to the FDA's labeling requirements that supports transparency to ensure that medical providers and the intended patient population has the information needed to use the device in a safe, effective, and appropriate manner for all.

AdvaMed Imaging thanks the Committee for the opportunity to submit this statement and hopes to serve as a partner and resource to the Committee. We hope to remain engaged with the Committee to ensure safe and appropriate access to ABHS.

AHIP

601 Pennsylvania Avenue, NW
South Building, Suite 500
Washington, DC 20004
T 202-778-3200
F 202-331-7487
<https://www.ahip.org/>

AHIP is the national association that represents health insurance providers, services, and solutions for millions of Americans. We are committed to market-based solutions and private-public partnerships while striving to enhance health care, both in terms of accessibility and affordability. We represent 128 health insurance plans nationwide. Collectively, our member plans provide access to health care for over 205 million people covered by employer-sponsored insurance, the individual insurance market, and public programs such as Medicare and Medicaid.

AHIP strives to be a valued and trusted resource for policymakers, regulators, and stakeholders that impact health care outcomes. As such, we welcome the opportunity to help inform policy discussions on artificial intelligence (AI) and its impact on health care.

As Americans increasingly encounter AI in every facet of life, including health care, it is important to create balanced policies that help realize the potential of AI and build trust among patients and stakeholders.¹ AI has the potential to meaningfully contribute to making health care more affordable, expanding access, and improving health outcomes. However, the promise of AI also comes with the potential for unintended consequences. As AI becomes further integrated into our health care systems, a robust and thoughtful policy approach will be crucial for advancing impactful applications and building trust among patients and stakeholders while preventing potential unintended consequences.

To that end, AHIP appreciates the Senate Committee on Finance's interest in the role of AI in health care. AHIP's members work every day to ensure that Americans have access to high-quality care and other supports, including appropriately using AI tools to fulfill the promise of guiding greater health. We look forward to working with the Committee and other stakeholders to enable a strong and resilient health system that deploys AI safely and responsibly, harnessing these technologies to drive quality, equitable, patient-centered, and affordable health care.

¹ <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC7325854/>.

The Role of Health Insurance Providers

AI has the potential to make the health care system work better and cost less. Health insurance providers are already using AI tools to reduce administrative costs, streamline and tailor consumer experience, improve and speed care, and minimize fraud.

Some examples of how health insurance providers use AI include:

- Analyzing provider directories to improve the accuracy of included data elements.
- Identifying patients who may benefit from improved access to services through predictive analytics.
- Assessing clinical performance to help consumers identify high-value care and for network design.
- Conducting claims analysis to reduce unnecessary spending and to identify potential fraud and abuse.
- Cleaning, normalizing, and labeling data for use in various programs.
- Conducting clinical models to understand health conditions and disease progression through research.
- Identifying gaps in the provision of evidenced-based care.
- Deploying service models to enhance the customer experience, such as chatbots.
- Streamlining prior authorization to identify data included in electronic medical records and speed requests and approvals (with clinicians involved throughout the process).
- Conducting actuarial analysis to help identify population-based utilization patterns for health plan sponsors and policymakers to understand usage trends both now and for the future.

We believe AI should augment, not replace, human decision-making and expertise and supplement existing clinical data to facilitate better decision-making. For example, AI can help facilitate a streamlined prior authorization process by using algorithms to issue approvals. In these tools, AI is generally only used to approve prior authorization requests for straightforward cases where the requisite documentation is provided by simply applying the same clinical algorithms that a person would use to approve a request, only faster. Denials are only automated for yes/no questions similar to whether a person is enrolled in the insurance plan, prior authorization is not required, or the service is categorically not included in the benefit. Cases requiring clinical decision-making for approvals or denials are individually reviewed by plan medical staff. The focused use of AI in prior authorizations allows health care providers and patients to receive approvals quickly, while health insurance providers can focus their experts on more complex cases.

The current advantages of AI highlight treatment innovations, care delivery transformation, operational efficiencies, cost reduction, technical innovation, and error reduction. AI has shown the true potential to both improve patient access to care and reduce administrative costs. While machine learning technologies have come a long way, the full extent of the benefits of AI applications have yet to be realized.

For instance, ambient and generative AI could take notes for physicians saving them time charting and allowing them to focus on patients. Natural language processing could find data in unstructured data in an electronic health record to reduce the burdens of quality measurement. Machine vision could analyze the gait of a recovering orthopedic patient at home to develop a physical therapy care plan remotely. Digital twins of consumers could allow providers and plans to continuously monitor enrollees' health. Clinical trials could be conducted primarily through simulations reducing ill effects on humans. There are endless possibilities if we strike the right balance between encouraging innovation and regulatory oversight.

Minimizing Bias

AHIP and its members are committed to ensuring that the application of AI is safe, transparent, explainable, and ethical. AHIP and our members seek to ensure these factors are integral components to AI systems, which will strengthen trust in the software techniques and outcomes.

We also seek to ensure biases are neither perpetuated nor introduced in the development and application of AI that could negatively impact certain subpopulations. While we may not be able to second guess and prevent all bias, robust monitoring and governance processes can enable swift course correction. Technology powered by AI can also play an important role in advancing health equity and improving health care access. For instance, health insurance providers can use predictive analytics to identify disparities in care and connect patients in need of additional services, such

as case management. Access to high-quality data sets, improving collection of demographic data, and leveraging industry consensus standards can support efforts to mitigate bias in AI.

Stakeholders in the private sector have been collaborating to develop governance, ethical, and practice standards for organizations developing and deploying AI to lead the way in protecting consumers while fostering AI. AHIP has joined business and technology leaders as well as consumer advocates to advance principles, best practices, and industry standards. For example, AHIP has worked with the Consumer Technology Association on developing standards of trustworthiness and recommendations for bias management.^{2,3}

AHIP is leading a broad-based multi-stakeholder effort to modernize and enhance demographic data content standards. Many overlapping and incomplete standards exist today. By seeking to ensure that standards are culturally sensitive, sufficiently granular, and aligned across stakeholders we seek to enable the collection of secure, accurate, complete, comparable, actionable, and interoperable data. In turn, access to these data will support better outcomes, fewer disparities, improved patient trust, and enhanced operational efficiency. Once the content standards are complete, exchange standards will be developed through the HL7 process. This will ensure we not only have the *what* but also the *where* (e.g., a questionnaire that can be built into an electronic health record or enrollment process) and the *how* (e.g., the ability to exchange the data between trusted partners). This information is critical to identifying disparities and determining where there may be bias in a model using AI so that bias can be mitigated and eliminated.

As these efforts evolve, AHIP encourages the Committee to foster public-private partnerships to invest in the necessary national infrastructure and consolidate and coalesce around common responsible standards on AI use.

Recommendations

As the use of AI grows, Americans deserve the peace of mind of knowing that it is being used responsibly and for their benefit. Thus, AHIP believes there is a need to set guardrails to protect consumers. However, those guardrails must reinforce policies that mitigate risks, including bias, without stifling innovation.

As the Committee considers AI oversight strategies, we urge you to consider ways to develop a streamlined and risk-based approach to legislation and oversight. We believe any legislation addressing AI should permit programs, practices, and procedures to reflect the context, scope, and data use of a specific use case. Legislation should allow for industry flexibility to tailor efforts to their unique circumstances in order to not impede innovation. AI legislation should focus on promoting appropriate governance, transparency, explainability, privacy, and mitigation techniques through adherence to industry and federal standards.

Implementing a Risk-Based Approach

Legislation should apply a risk-based approach to oversight that differentiates between “high-risk” or “high-impact” AI and low-risk AI. For example, the use of deep neural networks in clinical care presents more significant risks to patients than deploying simple algorithms to support administrative functions. Flexibility to right-size business practices and mitigation techniques based on risk is necessary to realize the potential of AI while avoiding overly restrictive, infeasible, or misaligned policies that risk stifling innovation.

AHIP strongly encourages the Committee to consider existing national frameworks and standards in developing any new legislation. This would help avoid duplication and ensure a streamlined oversight structure and support continued innovation. For example, the National Institute of Standards and Technology (NIST) AI Risk Management Framework, developed with robust input from stakeholders, including AHIP, provides a foundation for understanding and applying methods for “tiering” of risks associated with AI.⁴ The NIST AI Framework should be leveraged as Congress considers core components of AI legislation, such as definitions, to provide consistent direction to regulatory agencies in implementing any federal governance and oversight framework.

² <https://shop.cta.tech/collections/standards/products/the-use-of-artificial-intelligence-in-healthcare-trustworthiness-cta-2090>.

³ <https://shop.cta.tech/a/downloads/-/b5481e81fe7f99aa/9d1895627bdd6e27>.

⁴ <https://www.nist.gov/itl/ai-risk-management-framework>.

We further encourage the Committee to avoid subjecting *all* underlying AI technology to mandatory outside review or audit. Many health care organizations developing AI tools, particularly those who function as covered entities under HIPAA, are proactively employing their own risk-based approaches and optimizing existing data governance structures. Only large-scale foundation models, or general-purpose AI should be required to go through outside testing.

As AI will impact every industry, policymakers should take an “all-of-government” approach towards regulating these tools on a use-specific basis, fully leveraging existing industry standards and regulatory frameworks. This means balancing consistency across agencies with the need to tailor policies within the context of industry-specific use cases. For example, HHS, as the primary regulatory body for the health care sector, is best positioned to understand the intersection of any new AI oversight rules with existing regulatory frameworks, such as drug approvals or privacy requirements under HIPAA.

Fostering Transparency and Maintaining Public Trust

Trust is the foundation of our members’ engagement with patients and consumers. Health insurance providers build and maintain this trust today in numerous ways, including by protecting the privacy of patient information and promoting tools and resources to support patients’ active engagement in their health journey. Transparency is a key enabler of trust and is a critical component of successful deployment and use of AI. Patient, consumer, and caregiver education is critical to helping individuals better understand what AI is and how it might be used. AI transparency should also include information on what recourse individuals have if they believe it has been misused and resources to explain potential benefits (*e.g.*, clinical advancements) and potential drawbacks (*e.g.*, secondary uses of data). Americans will be more trusting of AI-based services with easy to access, understandable, and relevant information.

Understanding Accountability within the AI Ecosystem

Successful use of AI to enhance high-quality, equitable, and affordable care will depend on responsible approaches to both AI development and AI deployment. In the current environment, there are no established standards for delineating roles and responsibilities across the AI ecosystem, including between AI “developers” and “deployers,” if an adverse outcome occurs from the use of AI. There is also a lack of broadly applicable, established standards for transparency or disclosure of key elements of AI tools that would enable deployers to proactively assess the potential risk of errors or discrimination.⁵ We recommend that Congress look to established national frameworks, engage trusted federal partners, and leverage learnings from public-private efforts to inform any legislative policy efforts that address accountability in AI use.

Improving the Availability and Quality of Demographic Data

AI depends on its underlying data. For AI to function correctly, it is essential that the set or sets of data and other elements to calculate and achieve what has been programmed can be relied on as part of the function and methodology. Improving demographic data standards and collection will allow AI developers to build AI tools that benefit more people and prevent unintended bias that can result from not representing all populations in the data used to build and train AI. Congress should support the infrastructure necessary to capture data easily and consistently from patients and share it among trusted partners through policy changes and adoption of relevant standards by federal agencies such as the Office of Management and Budget, the Centers for Medicare & Medicaid Services, and the Office of the National Coordinator for Health Information Technology.

Leveraging Existing Laws and Regulations

As the Committee considers whether additional safeguards are necessary to protect against the potential risks of AI, such as algorithmic discrimination, we encourage you to consider how existing laws may already sufficiently protect people. Rather than develop numerous, conflicting laws and regulations, the federal government and the states should work together to leverage existing policies to foster transparency while preventing harm from AI.

⁵ The Office of the National Coordinator for Health IT (ONC) Health, Data, Technology, and Interoperability (HTI-1) final rule establishes transparency requirements for AI and other predictive algorithms that are part of certified health IT such as electronic health records. These requirements are applicable only to a narrow set of health care AI applications.

Ensuring Strong Privacy Protections

Everyone deserves the peace of mind of knowing that their personal health information is private, secure, and protected, regardless of who holds the data, or the type of technology application being used. As Congress considers privacy risks related to AI, AHIP recommends evaluating how current regulatory frameworks apply to mitigation of AI privacy risks. For example, HIPAA provides a robust framework to address privacy issues with respect to use of AI by HIPAA covered entities. HIPAA currently applies privacy and security rules to population-based programs to improve health or reduce health care costs. Across use cases, the same principles that exist in HIPAA today—such as minimum necessary, de-identification, notice, access, use and disclosure restrictions, and security requirements—should be applied to all entities holding health-related data regardless of whether AI is used in the process.

In addition to the extensive requirements of the HIPAA Privacy Rule, HIPAA covered entities must comply with a patchwork of state laws that are often conflicting or, in many cases, duplicative. We urge Congress to pass comprehensive national privacy legislation to fill this gap in our nation's privacy framework and avoid a 50-state patchwork quilt of rules. A coherent federal approach could help promote consistency in ensuring that the health data of all individuals is protected, wherever they reside or receive care now or in the future, and across a variety of services and technologies, to ensure no gaps or conflicts in privacy protection exist, including with respect to use of AI-enabled tools.

Congress should also seek to support the development of privacy enhancing technologies. AI can be used to further protect consumers through techniques such as pseudonymization, homomorphic encryption, secure multi-party computation, differential privacy, and zero knowledge proofs.

Conclusion

The appropriate use of AI holds great promise for improving health care for all Americans. AHIP believes that through public-private partnerships we can address the challenges posed by using AI while promoting innovation and maintaining American leadership. Engaging a diverse set of stakeholders is essential to this success. AHIP thanks the Committee for your attention to this critical issue. We look forward to working with you and other stakeholders on these important efforts.

AMERICAN FEDERATION OF TEACHERS

555 New Jersey Ave., NW
Washington, DC 20001
(202) 879-4400
<https://www.aft.org/>

February 20, 2024

United States Senate
Committee on Finance
Washington, DC 20510

Re: February 8th Finance Committee Hearing, Artificial Intelligence and Health Care: Promise and Pitfalls

Chairman Wyden:

On behalf of the AFT, the fastest-growing healthcare union in the nation, representing more than 200,000 healthcare professionals, as well as 1.72 million members who utilize the healthcare system, we are pleased to offer our thoughts on artificial intelligence for your hearing titled Artificial Intelligence and Health Care: Promise and Pitfalls.

The AFT believes that any discussion in this area must be centered around how artificial intelligence can benefit patients and healthcare workers. While there appears to be a useful role for AI in medical imaging and fighting cancer (and these positive uses may expand over time), AI cannot be used as a tool to replace healthcare professionals as a way to lower costs for hospitals and other institutions.

The Expertise of Professionals

I recently convened a meeting of healthcare leaders, and they expressed great concern about the impact of AI on patients. Although AI can be used to efficiently provide information to healthcare staff, there were concerns about using it for monitoring and diagnosis. My members offered examples of how having a physical presence in a hospital room is essential to properly evaluate, contextualize and deter-

mine patient care. The AFT supports legislation that ensures clinicians have the authority to exercise independent clinical judgment that deviates from AI-generated information, in the moment, at the bedside, without threat of disciplinary action or other retaliation. Further, workers should be at the table and have a voice at every phase of bringing AI into patient care—from the inception of the idea and discussions on how AI will be used, to testing and evaluating usage on an ongoing basis.

Healthcare Equity

There is a long history of bias in healthcare, and bias continues to have an impact ranging from disparities in maternal mortality to the development of medical devices like pulse oximeters. The AFT has concerns about the development and use of algorithms that may lead to biased outcomes. Your recent witness, Dr. Ziad Obermeyer, identified an algorithm making decisions using patients' predicted healthcare costs, rather than health needs, resulting in reinforced racial biases. The doctor also testified that "unfortunately, many of the biased algorithms we studied remain in use today."¹ The AFT is glad that the Food and Drug Administration has begun the process of seeking to regulate AI; and we hope that any recommendations include the frequent renewal of FDA authorization, as AI has been shown to potentially develop and expand on its own. More broadly, as included in the White House "Blueprint for an AI Bill of Rights,"² the AFT supports ongoing, independent evaluation of AI systems, as well as continued community and stakeholder input as systems develop to ensure they are assisting in undoing bias, not enforcing it.

Efficiency vs. Accuracy

Too many of our members have expressed concerns about being taken away from the bedside to have to complete documentation and deal with bureaucracy within the healthcare system. When AI can be used to reduce this paperwork burden in the clinical setting, or to improve note taking, it should be. Doing so will enhance and improve the delivery of patient care.

However, AI should not be used by payers as a tool to deny claims. Too many AFT members—whether they are educators, corrections officers, other public employees or healthcare professionals—struggle with healthcare affordability; and there have been enough examples of AI being used to deny coverage for our members to be concerned.³ The cases that have been brought to light are likely the tip of the iceberg, and patient care must always come before the profits of insurance companies. We are glad that the Centers for Medicare & Medicaid Services recently issued a notice to Medicare Advantage insurers providing guardrails on the use of AI and we hope that Congress will also take action to protect patients.⁴

Worker Protections

While our concerns, and this letter, are primarily focused on the needs of patients, as a leader of a labor union, I also want to highlight some important guardrails for employer-employee relations and AI. The integration of AI systems in the patient care setting, or for surveillance, should be recognized as an integral aspect of the working conditions and part of collective bargaining. In addition, any tracking or monitoring of healthcare workers should comply with all applicable privacy laws and regulations. Finally, when AI systems collect data about workers, that data should be minimized and not sold or commodified.

Thank you for organizing this hearing and considering our views on this important subject.

Sincerely,

Randi Weingarten
President

¹ Ziad Obermeyer, M.D., "AI Will Transform Medicine and the Health Care System—for Better or for Worse, Depending on How It Is Built and Applied." Senate Committee on Finance Hearing: Artificial Intelligence and Health Care: Promise and Pitfalls. Feb. 8, 2024. <https://www.finance.senate.gov/download/02042024-obermeyer-testimony>.

² Office of Science and Technology Policy, "Blueprint for an AI Bill of Rights." The White House, Oct. 4, 2022. <https://www.whitehouse.gov/ostp/ai-bill-of-rights/#discrimination>.

³ Lopez and Pugh, "AI Lawsuits Against Insurers Signal Wave of Health Litigation," *Bloomberg Law*, Feb. 1, 2024. <https://news.bloomberglaw.com/health-law-and-business/ai-law-suits-against-insurers-signal-wave-of-health-litigation>.

⁴ Bartnick, Cheng, Perumal, "CMS Confirms Medicare Advantage Organizations May Use AI in Making Coverage Determinations," Reed Smith Client Alerts, Feb. 13, 2024. <https://www.reedsmith.com/en/perspectives/2024/02/cms-confirms-medicare-advantage-organizations-may-use-ai-in-making-coverage>.

AMERICAN INSTITUTE FOR MEDICAL AND BIOLOGICAL ENGINEERING

1400 I St., NW, Suite 235
 Washington, DC 20005
 (202) 496-9660
<https://aimbe.org/>

AIMBE is a nonprofit, honorific organization representing the most accomplished leaders in the fields of medical and biological engineering across academia, industry, government, and scientific societies. AIMBE's mission is to provide advocacy in medical and biological engineering for the benefit of society.

AIMBE Fellows are at the forefront of health care innovation including developing artificial intelligence tools and algorithms for use in the clinic. Medical AI (also known as Health AI) has the potential to positively transform health care in the United States, but only to the extent its applications are deployed and utilized in clinical settings. Thus, investing in additional AI research, including open-access datasets, and incentivizing the use of Medical AI at the point of patient care is critical for tangible benefits to care and cost savings to be realized.

AI has several applications in medicine that can improve health care outcomes for patients through earlier detection, screening, and diagnosis. Research in the US has shown that by improving early diagnosis and personalizing treatment, AI can enhance the quality of medical care in terms of health outcomes and patient experience. For instance, lung cancer is associated with a 65% 5-year survival rate when it is localized compared to 9% when the disease has metastasized. By detecting disease earlier, AI models like Sybil that predict lung cancer risk can both save lives and significantly reduce overall cost associated with later-stage cancer treatments and care.¹ Moreover, in surgery settings, there is growing evidence that pre-operative cognitive state is a risk factor for postoperative adverse outcomes. Unfortunately, cognition is not assessed systematically pre- and postoperatively due to prohibitive costs and time constraints. AI tools have recently made it possible to quickly assess cognitive function in older adult surgical patients while significantly reducing nurse and administrative costs and overhead.² Use of tools such as this would greatly improve patient outcomes for the up to 65% of older patients that experience delirium and cognitive decline associated with surgical procedures each year.³

As the cost of medical treatment and health care in the United States continues to rise, cost-effective and value-based solutions are needed. According to CMS data, U.S. health care spending reached \$4.3 trillion or \$12,914 per person in 2021.⁴ As a share of the nation's Gross Domestic Product, health spending accounted for 18.3 percent. A recent study demonstrates that AI tools can provide tremendous cost savings in patient diagnosis and treatment. It is estimated that wider adoption of AI could lead to savings of 5 to 10% in U.S. health care spending—roughly \$200 billion to \$360 billion annually.⁵

Despite Medical AI providing cost-effective tools and being a rapidly growing area of biomedical research, its applications are severely underutilized in hospital and clinical health care settings. Even when Medical AI tools are available to clinicians and providers, they are disincentivized from using these tools without a reimbursement framework. While the US currently leads in the innovation of AI applications for medicine, it significantly lags behind the developing world in the adoption of its own tools. This gap will continue to widen without robust investment in AI research, large, open datasets, and prioritization of reimbursement pathways for AI. As a key funder of biomedical research in the world, our government has a duty to address critical bottlenecks between medical AI innovation and its use to directly improve patient health care.

Thank you in advance for considering the factors we have outlined in this statement. We appreciate the challenges and complexities you face as new AI tools are developed in the health sector. We stand ready to serve as a resource and assist your efforts as innovation continues and new policies and tools are needed.

¹<https://pubmed.ncbi.nlm.nih.gov/36634294/>.

²<https://pubmed.ncbi.nlm.nih.gov/37149670/>.

³<https://jamanetwork.com/journals/jama/fullarticle/2782851>.

⁴<https://www.cms.gov/data-research/statistics-trends-and-reports/national-health-expenditure-data/historical#:~:text=U.S.%20health%20care%20spending%20grew,For%20additional%20information%2C%20see%20below>.

⁵<https://www.nber.org/papers/w30857>.

AMERICAN MEDICAL ASSOCIATION
 25 Massachusetts Ave., NW, Suite 600
 Washington, DC 20001
 (202) 789-7400
<https://www.ama-assn.org/>

The American Medical Association (AMA) appreciates the opportunity to submit the following Statement for the Record to the U.S. Senate Committee on Finance as part of the hearing entitled, “Artificial Intelligence and Health Care: Promise and Pitfalls.” The AMA commends the Committee for its consideration of this critically important issue. Health care technology is advancing rapidly for many different uses and within many different sectors of the health care industry. Ensuring the responsible, equitable, ethical, and transparent design, development, and deployment of high-performing augmented intelligence (AI)-enabled tools within our health care system is a key priority for AMA members and our patients.¹ We strongly encourage the Committee to broadly ensure health care AI is considered as a sector of significant national concern and importance and to engage with health care stakeholders to ensure appropriate policies, standards, and regulatory requirements are in place to protect patient safety, promote equity, and ensure the quality and performance of the AI-enabled tools in question.

AMA Principles for Augmented Intelligence (AI) Development, Deployment, and Use

AMA’s new Principles for Augmented Intelligence (AI) Development, Deployment, and Use (<https://www.ama-assn.org/system/files/ama-ai-principles.pdf>), look to build on the AMA’s 2018 foundational principles on AI and seek to provide physician perspective on important AI policy topics. These principles represent the AMA’s next steps to provide guidance and drive change to help protect patients and physicians while recognizing the opportunities AI presents. We urge policy makers to move swiftly to implement AI guardrails that ensure the ethical, equitable, responsible, and transparent implementation of AI that both encourages safety and quality while promoting the opportunities presented by the emerging technologies.

In developing these principles, AMA members made clear that ensuring accurate performance and mitigating risks of AI are of the utmost priority. As such, these new principles urge Congress and the Administration to move more expediently towards developing national governance policies for implementation of AI-enabled technologies.

One key step in enhancing safety and limiting risk is to move decisively towards increased transparency requirements for health care AI. Physicians and patients must know when they are engaging with AI and they must be aware when medical decision-making includes consultation with AI-enabled technologies. Developers of AI-enabled technologies must disclose information about their products that allow purchasers and end users to fully evaluate the tool’s appropriateness, quality, performance, and risk of bias. We also must work swiftly to ensure AI includes strong data privacy protections and minimizes the ever-increasing cybersecurity risks continually plaguing hospitals, health systems and physician practices.

Additionally, we are growing increasingly concerned about the use of automated decision-making tools by health insurers, including many payors offering Medicare Advantage plans. Numerous reports of these tools increasing claims denials and limiting access to vital care show an urgent need to limit the use of these technologies in claims determinations that result in denials of care and limitations on coverage. While certain uses of AI-enabled technologies may increase efficiencies and reduce administrative burdens, claims determinations should still be reviewed on the health circumstances of the individual in question and should not depend on a standardized algorithm treating every patient the same. While we support the U.S. Department of Health and Human Services’ recent policies to curb AI use by Medicare Advantage organizations, more must be done to ensure health insurers comply with Medicare’s rules and do not create barriers to care. We strongly urge Congress to ensure that AI-enabled technologies and automated decision-making tools are used in limited and appropriate circumstances by insurers.

The AMA is pleased to see the growing focus on AI but is concerned these technologies will rapidly outpace regulations to ensure quality and safety if Congress

¹ The AMA refers to AI as “augmented intelligence”—a crucial concept in health care that emphasizes the enhancement of human decision-making through AI technologies, rather than replacing human expertise, ensuring a synergistic partnership where AI tools assist health care professionals in delivering more accurate, efficient, and personalized patient care.

and the Administration do not move quickly to ensure governance policies and guardrails are in place to guide implementation. Voluntary agreements among technology companies are simply not enough to give physicians and patients the assurances they need to pursue use of these tools. Congress and the Administration must work closely with not just big tech, but with physicians and patients, on appropriate policies to mitigate the potential risks of AI. We will only recognize the promise of these emerging technologies if we ensure that they are safe and if we can protect our patients from harm.

**American Medical Association
Principles for Augmented Intelligence Development, Deployment, and Use**

~ Approved by AMA Board of Trustees on November 14, 2023 ~

As the number of Augmented Intelligence (AI)-enabled health care tools and systems continue to grow, these technologies must be designed, developed, and deployed in a manner that is ethical, equitable, responsible, and transparent. With a lagging effort towards adoption of national governance policies or oversight of AI, it is critical that the physician community engage in development of policies to help inform physician and patient education, and guide engagement with these new technologies. It is also important that the physician community help guide development of these tools in a way that best meets both physician and patient needs, and help define their own organization's risk tolerance, particularly where AI impacts direct patient care. The AMA is committed to ensuring that AI can meet its full potential to advance clinical care and improve clinician well-being. This may only be accomplished by ensuring that physicians engage only with AI that satisfies rigorous standards to meet the goals of the quadruple aim,² advance health equity, prioritize patient safety, and limit risks to both physicians and patients.

These new principles build on earlier AMA policy development activities, including the 2018 foundational AMA AI policy, Augmented Intelligence in Medicine,³ followed by 2019 policy for payment and coverage of AI.⁴ However, as AI has rapidly developed beyond AI-enabled medical devices, new policy and guidance for adoption of both device and non-device uses of AI-enabled technologies is necessary to assist in deployment of these new advances to physicians and patients. These principles will serve as the foundation for AMA's evolving advocacy on AI.

The AMA is dedicated to providing continued guidance to physicians on how to best engage with new AI-enabled technologies with the understanding that policy development related to AI will likely continue to develop given the rapid pace of change in this space.

Oversight of Health Care Augmented Intelligence

There is currently no national policy or governance structure in place to guide the development and adoption of non-device AI. While the Food and Drug Administration (FDA) regulates AI-enabled medical devices, many types of AI-enabled technologies fall outside the scope of FDA oversight, including AI that may have clinical applications, such as some clinical decision support functions. While the Federal Trade Commission and the Health and Human Services Office for Civil Rights have oversight over some aspects of AI, their authorities are limited and not adequate to ensure appropriate development and deployment of AI generally, and specifically in the health care space. The AMA encourages a whole of government approach to implement governance policies that ensure overall and disparate risks to consumers and patients arising from AI are mitigated to the greatest extent possible.

In addition to government, health care institutions, practices, and professional societies share some responsibility for appropriate oversight and governance of AI-enabled systems and technologies. Beyond government oversight or regulation, purchasers and users of these technologies should have appropriate and sufficient policies in place to ensure they are acting in accordance with the current standard of care. Similarly, clinical experts are best positioned to determine whether AI applica-

²AI systems should enhance the patient experience of care and outcomes, improve population health, reduce overall costs for the health care system while increasing value, and support the professional satisfaction of physicians and the health care team.

³American Medical Association. (2018). "AI in Healthcare: A Report from the American Medical Association Board of Trustees." AMA, <https://www.ama-assn.org/system/files/2019-08/ai-2018-board-report.pdf> (Accessed September 14, 2023).

⁴American Medical Association. (2019). "AI in Healthcare: A Report from the American Medical Association Board of Trustees—2019." <https://www.ama-assn.org/system/files/2019-08/ai-2019-board-report.pdf> (Accessed September 14, 2023).

tions are high quality, appropriate, and whether the AI tools are valid from a clinical perspective. Clinical experts can best validate the clinical knowledge, clinical pathways, and standards of care used in the design of AI-enabled tools and can monitor the technology for clinical validity as it evolves over time.

- Health care AI must be designed, developed, and deployed in a manner which is ethical, equitable, responsible, and transparent.
- Use of AI in health care delivery requires clear national governance policies to regulate its adoption and utilization, ensuring patient safety, and mitigating inequities. Development of national governance policies should include inter-departmental and interagency collaboration.
- Compliance with national governance policies is necessary to develop AI in an ethical and responsible manner to ensure patient safety, quality, and continued access to care. Voluntary agreements or voluntary compliance is not sufficient.
- Health care AI requires a risk-based approach where the level of scrutiny, validation, and oversight should be proportionate to the potential overall or disparate harm and consequences the AI system might introduce. [See also Augmented Intelligence in Health Care H-480.939 at <https://policysearch.ama-assn.org/policyfinder/detail/H-480.939%20?uri=%2FAMADoc%2FHOD.xml-H-480.939.xml>.]
- Clinical decisions influenced by AI must be made with specified human intervention points during the decision-making process. As the potential for patient harm increases, the point in time when a physician should utilize their clinical judgment to interpret or act on an AI recommendation should occur earlier in the care plan.
- Health care practices and institutions should not utilize AI systems or technologies that introduce overall or disparate risk that is beyond their capabilities to mitigate. Implementation and utilization of AI should avoid exacerbating clinician burden and should be designed and deployed in harmony with the clinical workflow.
- Medical specialty societies, clinical experts, and informaticists are best positioned and should identify the most appropriate uses of AI-enabled technologies relevant to their clinical expertise and set the standard of care for AI usage in their specific domain. [See Augmented Intelligence in Health Care H-480.940 at <https://policysearch.ama-assn.org/policyfinder/detail/H-480.939%20?uri=%2FAMADoc%2FHOD.xml-H-480.939.xml>.]

When to Disclose: Transparency in Use of Augmented Intelligence-Enabled Systems and Technologies

As implementation of AI-enabled tools and systems continues to increase, it is essential that use of AI in health care be transparent to both physicians and patients. Transparency requirements should be tailored in a way that best suits the needs of the end users. Disclosure should contribute to physician and patient knowledge and not create unnecessary administrative burden. When AI is utilized in health care decision-making, that use should be disclosed and documented in order to limit risks to, and mitigate inequities for, both physicians and patients, and to allow each to understand how decisions impacting patient care or access to care are made. While transparency does not necessarily ensure AI-enabled tools are accurate, secure, or fair, it is difficult to establish trust if certain characteristics are hidden.

- When AI is used in a manner which directly impacts patient care, access to care, or medical decision making, that use of AI should be disclosed and documented to both physicians and/or patients in a culturally and linguistically appropriate manner. The opportunity for a patient or their caregiver to request additional review from a licensed clinician should be made available upon request.
- When AI is used in a manner which directly impacts patient care, access to care, medical decision making, or the medical record, that use of AI should be documented in the medical record.
- AI tools or systems cannot augment, create, or otherwise generate records, communications, or other content on behalf of a physician without that physician's consent and final review.
- When health care content is generated by generative AI, including by large language models, it should be clearly disclosed within the content that was generated by an AI-enabled technology.
- When AI or other algorithmic-based systems or programs are utilized in ways that impact patient access to care, such as by payors to make claims determinations or set coverage limitations, use of those systems or programs must be disclosed to impacted parties.

- The use of AI-enabled technologies by hospitals, health systems, physician practices, or other entities, where patients engage directly with AI should be clearly disclosed to patients at the beginning of the encounter or interaction with the AI-enabled technology.

What to Disclose: Required Disclosures by Health Care Augmented Intelligence-Enabled Systems and Technologies

Along with significant opportunity to improve patient care, all new technologies in health care will likely present certain risks and have limitations that physicians must carefully navigate during the early stages of clinical implementation of these new systems and tools. AI-enabled tools are no different and are perhaps more challenging than other advances as they present novel and complex questions and risks. To best mitigate these risks, it is critical that physicians understand AI-driven technologies and have access to certain information about the AI tool or system being considered, including how it was trained and validated, so that they can assess the quality, performance, equity, and utility of the tool to the best of their ability. This information may also establish a set of baseline metrics for comparing AI tools. Transparency and explainability regarding the design, development, and deployment processes should be mandated by law where possible, including potential sources of inequity in problem formulation, inputs, and implementation. Additionally, sufficient detail should be disclosed to allow physicians to determine whether a given AI-enabled tool would reasonably apply to the individual patient they are treating. Physicians should understand that, where they utilize AI-enabled tools and systems without transparency provided by the AI developer, their risks of liability for reliance on that AI will likely increase. The need for full transparency is greatest where AI-enabled systems have greater impacts on direct patient care, such as by AI-enabled medical devices, clinical decision support, and interaction with AI-driven chatbots. Transparency needs may be somewhat lower where AI is utilized for primarily administrative, practice-management functions.

- When AI-enabled systems and technologies are utilized in health care, the following information should be disclosed by the AI developer to allow the purchaser and/or user (physician) to appropriately evaluate the system or technology prior to purchase or utilization:
 - Regulatory approval status
 - Applicable consensus standards and clinical guidelines utilized in design, development, deployment, and continued use of the technology
 - Clear description of problem formulation and intended use accompanied by clear and detailed instructions for use
 - Intended population and intended practice setting
 - Clear description of any limitations or risks for use, including possible disparate impact
 - Description of how impacted populations were engaged during the AI lifecycle
 - Detailed information regarding data used to train the model:
 - Data provenance
 - Data size and completeness
 - Data timeframes
 - Data diversity
 - Data labeling accuracy
 - Validation Data/Information and evidence of:
 - Clinical expert validation in intended population and practice setting and intended clinical outcomes
 - Constraint to evidence-based outcomes and mitigation of “hallucination” or other output error
 - Algorithmic validation
 - External validation processes for ongoing evaluation of the model performance, *e.g.*, accounting for AI model drift and degradation
 - Comprehensiveness of data and steps taken to mitigate biased outcomes
 - Other relevant performance characteristics, including but not limited to performance characteristics at peer institutions/similar practice settings
 - Post-market surveillance activities aimed at ensuring continued safety, performance, and equity
 - Data Use Policy
 - Privacy
 - Security

- Special considerations for protected populations or groups put at increased risk
 - Information regarding maintenance of the algorithm, including any use of active patient data for ongoing training
 - Disclosures regarding the composition of design and development team, including diversity and conflicts of interest, and points of physician involvement and review
- Physicians should carefully consider whether or not to engage with AI-enabled health care technologies if this information is not disclosed by the developer. As the risk of AI being incorrect increases risks to patients (such as with clinical applications of AI that impact medical decision making), disclosure of this information becomes increasingly important. [See also Augmented Intelligence in Health Care H-480.939 <https://policysearch.ama-assn.org/policyfinder/detail/H-480.939?uri=%2FAMADoc%2FHOD.xml-H-480.939.xml>.]

Generative Augmented Intelligence

Generative AI is a type of AI that can recognize, summarize, translate, predict, and generate text and other content based on knowledge gained from large datasets. Generative AI tools are finding an increasing number of uses in health care, including assistance with administrative functions, such as generating office notes, responding to documentation requests, and generating patient messages. Additionally, there has been increasing discussion about clinical applications of generative AI, including use as clinical decision support to provide differential diagnoses, early detection, and intervention, and to assist in treatment planning. While generative AI tools show tremendous promise to make a significant contribution to health care, there are a number of potential risks and limitations to consider when using these tools in a clinical setting or direct patient care. To manage risk, health care organizations should develop and adopt appropriate policies that anticipate and minimize negative impacts. Physicians who are considering utilizing a generative AI-based tool in their practice should ensure that all practice staff are educated on the risks and limitations, including patient privacy concerns, and additionally, should have appropriate governance policies in place for its use prior to adoption.

- Generative AI should: (a) only be used where appropriate policies are in place within the practice or other health care organization to govern its use and help mitigate associated risks; and (b) follow applicable state and federal laws and regulations (*e.g.*, HIPAA-compliant Business Associate Agreement).
- Appropriate governance policies should be developed by health care organizations and account for and mitigate risks of:
 - Incorrect or falsified responses; lack of ability to readily verify the accuracy of responses or the sources used to generate the response
 - Training data set limitations that could result in responses that are out of date or otherwise incomplete or inaccurate for all patients or specific populations
 - Lack of regulatory or clinical oversight to ensure performance of the tool
 - Bias, discrimination, promotion of stereotypes, and disparate impacts on access or outcomes
 - Data privacy
 - Cybersecurity
 - Physician liability associated with the use of generative AI tools
- Health care organizations should work with their AI and other health information technology (health IT) system developers to implement rigorous data validation and verification protocols to ensure that only accurate, comprehensive, and bias managed datasets inform generative AI models, thereby safeguarding equitable patient care and medical outcomes. [See Augmented Intelligence in Health Care H-480.940 at <https://policysearch.ama-assn.org/policyfinder/detail/H-480.940?uri=%2FAMADoc%2FHOD.xml-H-480.940.xml>.]
- Use of generative AI should incorporate physician and staff education about the appropriate use, risks, and benefits of engaging with generative AI. Additionally, physicians should engage with generative AI tools only when adequate information regarding the product is provided to physicians and other users by the developers of those tools.
- Clinicians should be aware of the risks of patients engaging with generative AI products that produce inaccurate or harmful medical information (*e.g.*, patients asking chatbots about symptoms) and should be prepared to counsel patients on the limitations of AI-driven medical advice.

- Governance policies should prohibit the use of confidential, regulated, or proprietary information as prompts for generative AI to generate content.
- Data and prompts contributed by users should primarily be used by developers to improve the user experience and AI tool quality and not simply increase the AI tool's market value or revenue generating potential.

Physician Liability for Use of Augmented Intelligence-Enabled Technologies

The question of physician liability for use of AI-enabled technologies presents novel and complex legal questions and potentially poses risks to the successful clinical integration of AI-enabled technologies. As legal theories of liability and accountability for AI continue to evolve, the AMA will continue to advocate to ensure that physician liability for the use of AI-enabled technologies is limited and adheres to current legal approaches to medical malpractice.

- Current AMA policy states that liability and incentives should be aligned so that the individual(s) or entity(ies) best positioned to know the AI system risks and best positioned to avert or mitigate harm do so through design, development, validation, and implementation. [See Augmented Intelligence in Health Care H-480.939 <https://policysearch.ama-assn.org/policyfinder/detail/H-480.939%20?uri=%2FAMADoc%2FHOD.xml-H-480.939.xml>.]
- Where a mandated use of AI systems prevents mitigation of risk and harm, the individual or entity issuing the mandate must be assigned all applicable liability.
- Developers of autonomous AI systems with clinical applications (screening, diagnosis, treatment) are in the best position to manage issues of liability arising directly from system failure or misdiagnosis and must accept this liability with measures such as maintaining appropriate medical liability insurance and in their agreements with users.
- Health care AI systems that are subject to non-disclosure agreements concerning flaws, malfunctions, or patient harm (referred to as gag clauses) must not be covered or paid and the party initiating or enforcing the gag clause assumes liability for any harm.
- When physicians do not know or have reason to know that there are concerns about the quality and safety of an AI-enabled technology, they should not be held liable for the performance of the technology in question.

Data Privacy and Augmented Intelligence

Data privacy is highly relevant to AI development, implementation, and use. The AMA is deeply invested in ensuring individual patient rights and protections from discrimination remain intact, that these assurances are guaranteed, and that the responsibility falls with the data holders. AI development, training, and use requires assembling large collections of health data. AI machine learning is data hungry; it requires massive amounts of data to function properly. Increasingly, more electronic health records are interoperable across the health care system and, therefore, are accessible by AI trained or deployed in medical settings. AI developers may create legal arrangements, *e.g.*, business associate agreements, that bring them under the Health Insurance Portability and Accountability Act (HIPAA) Privacy and Security Rules. Yet even HIPAA cannot protect patients from the “black box” nature of AI which makes the use of data opaque. AI system outputs may also include inferences that reveal personal data or previously confidential details about individuals. This can result in a lack of accountability and trust and exacerbate data privacy concerns. Often, AI developers and implementers are themselves unaware of exactly how their products use information to make recommendations.

It is unlikely that physicians or patients will have any clear insight into a generative AI tool's conformance to state or federal data privacy laws. Large language models (LLM) are trained on data scraped from the web and other digital sources (including HIPAA-covered environments),⁵ yet few, if any, controls are available to help users protect the data they voluntarily enter in a chatbot query. For instance, there are often no mechanisms in place for users to request data deletion or ensure that their inputs are not stored or used for future model training. While tools designed for medical use should align with HIPAA, many “HIPAA-compliant” generative tools rely on antiquated notions of deidentification, *i.e.*, stripping data of per-

⁵Feathers, T., et al. “Facebook is receiving sensitive medical information from hospital websites. The Markup. June 16, 2022.” <https://themarkup.org/pixel-hunt/2022/06/16/facebook-is-receiving-sensitive-medical-information-from-hospital-websites>.

sonal information. With today's advances in computing power, data can easily be re-identified. Rather than aiming to make LLMs compliant with HIPAA, all health care AI-powered generative tools should be designed from the ground up with data privacy in mind.

The AMA's Privacy Principles (<https://www.ama-assn.org/system/files/2020-05/privacy-principles.pdf>) were designed to provide individuals with rights and protections and shift the responsibility for privacy to third-party data holders. While the Principles are broadly applicable to all AI developers, *e.g.*, entities should only collect the minimum amount of information needed for a particular purpose, the unique nature of LLMs and generative AI warrant special emphasis on entity responsibility and user education.

Entity Responsibility:

- Entities should make information available about the intended use of generative AI in health care and identify the purpose of its use. Individuals should know how their data will be used or reused, and the potential risks and benefits.
- Individuals should have the right to opt-out, update, or forget use of their data in generative AI tools. These rights should encompass AI training data and disclosure to other users of the tool.
- Generative AI tools should not reverse engineer, reconstruct, or reidentify an individual's originally identifiable data or use identifiable data for nonpermitted uses, *e.g.*, when data are permitted to conduct quality and safety evaluations. Preventive measures should include both legal frameworks and data model protections, *e.g.*, secure enclaves, federated learning, and differential privacy.

User Education:

- Users should be provided with training specifically on generative AI. Education should address:
 - legal, ethical, and equity considerations,
 - risks such as data breaches and re-identification,
 - potential pitfalls of inputting sensitive and personal data, and
 - the importance of transparency with patients regarding the use of generative AI and their data.

[See Augmented Intelligence in Health Care H-480.940 at <https://policysearch.ama-assn.org/policyfinder/detail/H-480.940?uri=%2FAMADoc%2FHOD.xml-H-480.940.xml>.]

Augmented Intelligence Cybersecurity

Data privacy relies on strong data security measures. There is growing concern that cyber criminals will use AI to attack health care organizations. AI poses new threats to health IT operations. AI-operated ransomware and AI-operated malware can be targeted to infiltrate health IT systems and automatically exploit vulnerabilities. Attackers using ChatGPT can craft convincing or authentic emails and use phishing techniques that entice people to click on links—giving them access to the entire electronic health record system.

AI is particularly sensitive to the quality of data. Data poisoning is the introduction of “bad” data into an AI training set, affecting the model's output. AI requires large sets of data to build logic and patterns used in clinical decision-making. Protecting this source data is critical. Threat actors could also introduce input data that compromises the overall function of the AI tool. Failure to secure and validate these inputs, and corresponding data, can contaminate AI models—resulting in patient harm.

Because stringent privacy protections and higher data quality standards might slow model development, there could be a tendency to forgo essential data privacy and security precautions. However, strengthening AI systems against cybersecurity threats is crucial to their reliability, resiliency, and safety.

AI cybersecurity considerations:

- AI systems must have strong protections against input manipulation and malicious attacks.
- Entities developing or deploying health care AI should regularly monitor for anomalies or performance deviations, comparing AI outputs against known and normal behavior.
- Independent of an entity's legal responsibility to notify a health care provider or organization of a data breach, that entity should also act diligently in identi-

fying and notifying the individuals themselves of breaches that impact their personal information.

- Users should be provided education on AI cybersecurity fundamentals, including specific cybersecurity risks that AI systems can face, evolving tactics of AI cyber attackers, and the user's role in mitigating threats and reporting suspicious AI behavior or outputs.

Payor Use of Augmented Intelligence and Automated Decision-Making Systems

Payors and health plans are increasingly using AI and algorithm-based decision-making in an automated fashion to determine coverage limits, make claim determinations, and engage in benefit design. Payors should leverage automated decision-making systems that improve or enhance efficiencies in coverage and payment automation, facilitate administrative simplification, and reduce workflow burdens. While the use of these systems can create efficiencies such as speeding up prior authorization and cutting down on paperwork, there is concern these systems are not being designed or supervised effectively, creating access barriers for patients and limiting essential benefits.

Increasingly, evidence shows that payors are using automated decision-making systems to deny care more rapidly, often with little or no human review. This manifests in the form of increased denials, stricter coverage limitations, and constrained benefit offerings. For example, a payor allowed an automated system to cut off insurance payments for Medicare Advantage patients struggling to recover from severe diseases, forcing them to forgo care or pay out of pocket. In some instances, payors instantly reject claims on medical grounds without opening or reviewing the patient's medical record. There is also a lack of transparency in the development of automated decision-making systems. Rather than payors making determinations based on individualized patient care needs, reports show that decisions are based on algorithms developed using average or "similar patients" pulled from a database. Models that rely on generalized, historical data can also perpetuate biases leading to discriminatory practices or less inclusive coverage.^{6, 7, 8, 9}

We must ensure that automated decision-making systems do not reduce needed care, nor systematically withhold care from specific groups. Steps should be taken to ensure that these systems are not overriding clinical judgement. Patients and physicians should be informed and empowered to question a payor's automated decision-making. There should be stronger regulatory oversight, transparency, and audits when payors use these systems for coverage, claim determinations, and benefit design. [See Use of Augmented Intelligence for Prior Authorization D-480.956 <https://policysearch.ama-assn.org/policyfinder/detail/D-480.956?uri=%2FAMADoc%2Fdirectives.xml-D-480.956.xml>; Prior Authorization and Utilization Management Reform H-320.939, <https://policysearch.ama-assn.org/policyfinder/detail/H-480.939%20?uri=%2FAMADoc%2FHOD.xml-H-480.939.xml>.]

- Use of automated decision-making systems that determine coverage limits, make claim determinations, and engage in benefit design should be publicly reported, based on easily accessible evidence-based clinical guidelines (as opposed to proprietary payor criteria), and disclosed to both patients and their physician in a way that is easy to understand.
- Payors should only use automated decision-making systems to improve or enhance efficiencies in coverage and payment automation, facilitate administrative simplification, and reduce workflow burdens. Automated decision-making systems should never create or exacerbate overall or disparate access barriers to needed benefits by increasing denials, coverage limitations, or limiting benefit offerings. Use of automated decision-making systems should not replace the individualized assessment of a patient's specific medical and social circumstances and payors' use of such systems should allow for flexibility to override auto-

⁶Obermeyer, Ziad, et al. "Dissecting racial bias in an algorithm used to manage the health of populations." *Science* 366.6464 (2019): 447–453. <https://www.science.org/doi/10.1126/science.aax2342>.

⁷Ross, C., Herman, B. (2023) "Medicare Advantage Plans' Use of Artificial Intelligence Leads to More Denials." <https://www.statnews.com/2023/03/13/medicare-advantage-plans-denial-artificial-intelligence/> (Accessed September 14, 2023).

⁸Rucker, P., Miller, M., Armstrong, D. (2023). "Cigna and Its Algorithm Deny Some Claims for Genetic Testing, ProPublica Finds." <https://www.propublica.org/article/cigna-pxdx-medical-health-insurance-rejection-claims> (Accessed September 14, 2023).

⁹Ross, C., Herman, B. (2023). "Medicare Advantage Algorithms Lead to Coverage Denials, With Big Implications for Patients." <https://www.statnews.com/2023/07/11/medicare-advantage-algorithm-navihealth-unitedhealth-insurance-coverage/> (Accessed September 14, 2023).

mated decisions. Payors should always make determinations based on particular patient care needs and not base decisions on algorithms developed on “similar” or “like” patients.

- Payors using automated decision-making systems should disclose information about any algorithm training and reference data, including where data were sourced and attributes about individuals contained within the training data set (*e.g.*, age, race, gender). Payors should provide clear evidence that their systems do not discriminate, increase inequities, and that protections are in place to mitigate bias.
- Payors using automated decision-making systems should identify and cite peer-reviewed studies assessing the system’s accuracy measured against the outcomes of patients and the validity of the system’s predictions.
- Any automated decision-making system recommendation that indicates limitations or denials of care, at both the initial review and appeal levels, should be automatically referred for review to a physician (a) possessing a current and valid non-restricted license to practice medicine in the state in which the proposed services would be provided if authorized and (b) be of the same specialty as the physician who typically manages the medical condition or disease or provides the health care service involved in the request prior to issuance of any final determination. Prior to issuing an adverse determination, the treating physician must have the opportunity to discuss the medical necessity of the care directly with the physician who will be responsible for determining if the care is authorized.
- Individuals impacted by a payor’s automated decision-making system, including patients and their physicians, must have access to all relevant information (including the coverage criteria, results that led to the coverage determination, and clinical guidelines used).
- Payors using automated decision-making systems should be required to engage in regular system audits to ensure use of the system is not increasing overall or disparate claims denials or coverage limitations, or otherwise decreasing access to care. Payors using automated decision-making systems should make statistics regarding systems’ approval, denial, and appeal rates available on their website (or another publicly available website) in a readily accessible format with patient population demographics to report and contextualize equity implications of automated decisions. Insurance regulators should consider requiring reporting of payor use of automated decision-making systems so that they can be monitored for negative and disparate impacts on access to care. Payor use of automated decision-making systems must conform to all relevant state and federal laws.

We appreciate the Committee’s focus on this important and promising tool that has the potential to significantly enhance the patient-physician relationship and improve and streamline the health care delivery system. As our comments above suggest, there are many considerations that must be thoroughly taken into account to ensure that AI delivers on the promise this technology holds to improve our nation’s health care. National guidelines with stakeholder buy-in that include physician and end-user input is a great place to start. The AMA and the physician community stand ready to serve as a resource to help this Committee and the other Committees of jurisdiction embark on this effort.

ASHER INFORMATICS PBC

6401 Penn Ave., 3rd Floor
Pittsburgh, PA 15206
<https://www.asherinformatics.com/>

Statement of John F. Kalafut, Ph.D., Chief Strategy Officer

Asher Informatics PBC is heartened and encouraged by the hearing on “Artificial intelligence and Healthcare: Promise & Pitfalls” convened by Chairman Ron Wyden (D-Oregon) and the Senate Finance Committee. Their staff assembled a stellar ensemble of witnesses that, we think, brought credible and practical information to their testimony. The witnesses represent the varied and important constituencies required to develop and maintain credible and equitable policies to responsibly catalyze the deployment and use of data-driven and algorithm-based interventions in clinical care. While the use of computer-aided decision tools and “AI” are not new to some areas of medicine (*e.g.*: radiology, anesthesia, neurology, audiology), the recent convergence of computational power, curated datasets, and breakthroughs in algorithm design allows us to contemplate the use of algorithm-based health serv-

ices (per Peter Shen at Siemens) across the spectrum of healthcare delivery—the back-office to the bedside. This is a nuanced and complex domain, however. Creating meaningful health outcomes with the new generation of AI tools is not simply allowing ChatGPT to be used across a hospital. “Generative AI” and its most recognizable interaction model—as a “chatbot”—is not the only type of AI nor is it usually the most appropriate type for use in high-risk and mission-critical scenarios as in clinical care decisions and treatment planning. Understanding the landscape of algorithm-based offerings for healthcare requires expertise and insight that already overburdened health systems typically don’t have.

The “FOMO” caused by the over-exuberant promotion of “genAI” by technology evangelists can potentially lead to health-systems making ill-informed decisions about AI adoption, especially if the decisions are shaped by large technology firms who are incentivized to sell large amount of storage and compute. As Dr. Mark Sendak points out in his testimony, if large and well-resourced health-systems are the ones most likely to have requisite expertise and infrastructure to deploy and use clinical AI effectively, we will again be risking the widening of health-inequities. CMS (and private payers) can help “level the playing field,” though, as suggested by Peter Shen of Siemens Healthineers through more uniform payment “boosts” or “add-on” payments to help accelerate the adoption and use of AI applications. This is particularly relevant for the broad class of AI applications on the market that have received pre-market clearance by the FDA (granted, Siemens Healthineers sells many of these types of products).

Doing so would not be without precedent, either. In the late 1990s and early 2000s, CMS reimbursement mechanisms that enabled providers to receive small, add-on payments if Computer Aided Detection (CAD) software was used in the interpretation and reporting of screening mammography. Later in the decade, procedural codes and reimbursement formulae were also adjusted to allow breast imagers to get small reimbursement boosts for using computer vision/processing CAD tools for interpreting MRI of the breast. There also were professional reimbursement boosts developed for radiologists to use advanced visualization software when interpreting some studies. Most of these reimbursements have either ended, been dramatically reduced, or will be authorized by private payers in very specific encounters. The first generation of Breast and Chest CAD did ultimately disappoint when larger, effectiveness studies were done and published.

There are multiple reasons for the sub-optimal results and eventual dissatisfaction of breast imagers with CAD—technical, usability/human-factors, data policy, and interoperability flaws. But the reimbursement drove the diffusion and use of breast CAD into the clinic which also allowed for the broader assessment and measurement of the technology’s utility. The first CAD product for mammography was approved by the FDA in 1998 and the US CMS issued breast CAD reimbursement codes (with RVUS) in 2002. By 2008, 74% of all breast mammograms read in the US were assisted with the CAD software, increasing to 92% by 2016 [Gao et al. reference AJR new frontiers in breast AI/CAD].

One could mistakenly view the experience of first-generation “AI” in medicine as a lesson in wastefulness driven by non-usable technology and therefore any computer-aided intervention or algorithmically assisted technology should be rejected outright or be saddled with evidence burdens that will stifle innovation because promising methods will require clinical evidence and testing that will far outstrip the resources of small and medium sized enterprises.

The landscape today is also quite different from the first-generation of AI tools in diagnostic medicine; the computational methods are far superior, it is easier (though still hard) to aggregate large medical datasets, there are more modern paradigms of supporting computational applications in health systems, and clinical data are digital and somewhat exchangeable across venues for research and clinical utility measurement. We think the example of 1st Gen AI is an example of the system working! It’s not as if the CAD technologies didn’t have any positive effects or lacked any evidence necessary to allow their diffusion. Most healthcare technologies need wider study and assessment to arrive at a determination of effectiveness.

To make meaningful improvements to healthcare via technology, we can’t expect every developer to have the budget or resources of large pharma. Similarly, not every health innovation should or requires Randomized Control Trials to demonstrate utility. There is an old joke in the bio-statistics community similar to: “the RCT for the new parachute system was stopped prematurely after the first control-arm subject was tested”. There tends to be in certain quarters of medicine a zealot-like assertion that we shouldn’t adopt any new method unless it can be run through

an RCT. Yes—we need to strive towards strong and convincing evidence of all medical interventions, but sometimes common-sense, realistic constraints, or “good enough is good enough” need to be considered. Did we really need RCTs to demonstrate that CT scanning would cause a seismic shift in patient care? Should there have been thousands of sham or real surgeries to open-up patients and compare to treatment costs for patients with just images?

CENTER FOR AI AND DIGITAL POLICY

Open Gov Hub
1100 13th St., NW, Suite 800
Washington, DC 20005

February 14, 2024

Chair Ron Wyden
Ranking Member Mike Crapo
U.S. Senate Committee on Finance
219 Dirksen Senate Office Building
Washington, DC 20510

Re: CAIDP Statement for the Record: “Artificial Intelligence and Health Care: Promise and Pitfalls”

Dear Chairman Wyden, Ranking Member Crapo, and Members of the Committee,

We write to you regarding the hearing on “Artificial Intelligence and Health Care: Promise and Pitfalls.”¹ The Center for AI and Digital Policy (CAIDP) appreciates your leadership on addressing the risks and benefits of AI systems and your work towards establishing standards of AI governance.

The CAIDP is an independent research and education non-profit based in Washington, DC.² Our global network of AI policy experts and advocates advises national governments, international organizations, and congressional committees regarding artificial intelligence and digital policy. Our President, Merve Hickok testified at the first congressional hearing on AI last year—“Advances in AI: Are We Ready for a Tech Revolution?”³ CAIDP routinely provides advice to Congressional Committees on matters involving AI policy. We previously advised the Senate Judiciary Committee on AI in Criminal Prosecutions,⁴ AI and Human Rights,⁵ the Senate HELP Committee on AI and Healthcare,⁶ and the Senate Rules Committee on AI and Elections.⁷

We also publish the annual Artificial Intelligence and Democratic Values Report,⁸ providing a comprehensive review of AI policies and practices in 75 countries.

In brief, our recommendations to this Committee are:

- 1) Exercise oversight of federal agencies tasked with ensuring the responsible use of AI in the health care sector under the Biden AI Executive Order.⁹

¹ U.S. Senate Committee on Finance, *Artificial Intelligence and Health Care: Promise and Pitfalls* (Feb. 8, 2024), <https://www.finance.senate.gov/hearings/artificial-intelligence-and-health-care-promise-and-pitfalls>.

² CAIDP, About, <https://www.caidp.org/about-2/>.

³ Testimony and statement for the record of CAIDP President Merve Hickok, *Advances in AI: Are We Ready For a Tech Revolution?*, House Committee on Oversight and Accountability, Subcommittee on Cybersecurity, Information Technology, and Government Innovation (March 8, 2023), https://oversight.house.gov/wp-content/uploads/2023/03/Merve-Hickok_testimony_March-8th-2023.pdf.

⁴ CAIDP, *Statements*, <https://www.caidp.org/statements/>.

⁵ CAIDP, *Statement to Senate Judiciary Committee on “AI and Human Rights”* (June 13, 2023), <https://www.caidp.org/app/download/8462575863/CAIDP-SJC-06132023.pdf>.

⁶ CAIDP, *Statement to Senate HELP Committee on “Avoiding a Cautionary Tale: Policy Considerations for Artificial Intelligence in Healthcare”* (Nov. 8, 2023), <https://www.caidp.org/app/download/8487454163/CAIDP-Senate%20HELP-AI-Healthcare-11082023.pdf>.

⁷ CAIDP, *Statement to Senate Rules Committee on “AI and Elections”* (September 27, 2023), <https://www.caidp.org/app/download/8478562663/CAIDP-SRC-AI-ELECTIONS-09272023.pdf>.

⁸ CAIDP, *Artificial Intelligence and Democratic Values* (2023), <https://www.caidp.org/reports/aidv-2023/>.

⁹ Executive Order 14110 of October 30, 2023, *Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*, Federal Register Vol. 88, No. 210, pg. 75191–75226, <https://www.govinfo.gov/content/pkg/FR-2023-11-01/pdf/2023-24283.pdf>.

- 2) Move forward AI legislation. We endorse the Algorithmic Accountability Act of 2023 and the Blumenthal-Hawley Framework for a U.S. AI Act.

AI and Healthcare

The Food and Drug Administration has allowed medical algorithms since 1995, mostly for medical imaging.¹⁰ The promise of healthcare AI has been in efforts to improve drug development, detect diseases earlier, and analyze medical data more consistently.¹¹ A Wired report¹² highlights extensive bias and discrimination produced by predictive systems deployed in mental and physical healthcare, impacting non-clinical and diagnostic services. “Female patients are disproportionately misdiagnosed for heart disease and receive insufficient or incorrect treatment.”¹³ Another study evidences significant racial bias in a widely used pricing algorithm, affecting millions of patients¹⁴ and suggests that remedying this disparity would increase the percentage of Black patients receiving proper medical care from 17.7 to 46.5%.¹⁵ The bias arises because the algorithm predicts health care costs rather than illness, but unequal access to care means that we spend less money caring for Black patients than for White patients.

The application of AI/ML systems in generating consumer reports and insurance scoring decisions pose particular risks for ensuring fair and equitable access to healthcare for Americans. Regarding algorithmic decision systems in healthcare, CAIDP President Merve Hickok conducted a study that highlights the healthcare sector’s wealth of data, emphasizing the need for safe and ethical development, deployment, and implementation of algorithmic tools. Failure to do so could have harmful effects on patients’ lives, well-being, and safety.¹⁶

The Congressional Research Service Report highlights concerns regarding the accuracy, security, and privacy of AI technologies in healthcare. These concerns include the availability of sufficient health data, medical liability, consent processes, and patient access. Hence, risks associated with AI systems also include misdiagnosis due to poor design, amplification of biases in data, and potential widespread patient injury if flawed systems are widely adopted.¹⁷

The use of AI in diagnostic services, insurance, and medical coverage has led to several disputes before the courts on fraud¹⁸ and abusive business practices,¹⁹ exposing real risks to the public. Moreover, the expansion of AI into nonclinical areas of healthcare raises critical considerations for patient safety and efficacy. While AI algorithms do not require FDA clearance if they do not directly impact clinical

¹⁰ HealthExec, *FDA has now cleared more than 500 healthcare AI algorithms* (Feb. 6, 2023), <https://healthexec.com/topics/artificial-intelligence/fda-has-now-cleared-more-500-healthcare-ai-algorithms> [“HealthExec Report”].

¹¹ Congressional Research Service, *Artificial Intelligence: Overview, Recent Advances, and Considerations for the 118th Congress* (Aug. 4, 2023), pg. 4, <https://crsreports.congress.gov/product/pdf/R/R47644> [“CRS Report: 118th Congress”].

¹² Wired, *Health Care Bias Is Dangerous. But So Are “Fairness” Algorithms* (Feb. 8, 2023), <https://www.wired.com/story/bias-statistics-artificial-intelligence-healthcare/>.

¹³ *Id.*

¹⁴ Ziad Obermeyer et al., *Dissecting racial bias in an algorithm used to manage the health of populations*, *Science*, 366,447–453 (2019), DOI:10.1126/science.aax2342, <https://www.science.org/doi/10.1126/science.aax2342>.

¹⁵ *Id.*

¹⁶ Merve Hickok, Colleen Dorsey, Tim O’Brien, Dorothea Baur, Katrina Ingram, Chhavi Chauhan, Attlee M. Gamundani, *Case Study: The Distilling of a Biased Algorithmic Decision System through a Business Lens*, DOI: 10.31235/osf.io/t5dhu; <https://osf.io/preprints/socarxiv/t5dhu/>.

¹⁷ CRS Report: 118th Congress, pg. 5.

¹⁸ *Braun v. Ontrak, Inc.*, 2023 Cal. Super. LEXIS 71440. [Ontrak-A program was that the insured patients targeted for recruiting into the program were also disproportionately people who were more likely to lose their health coverage due to job loss or other causes. As a result, although Ontrak would represent to patients that they would not be required to pay for Ontrak’s services, by the time Ontrak billed its insurer customers for the services provided to the patients, the clients were often no longer covered by their insurance.]

¹⁹ *In Re Meta Pixel Healthcare Litigation*, 647 F. Supp. 3d 778, 784. [Plaintiffs are Facebook users who allege that Met improperly acquires their confidential health information in violation of state and federal law and in contravention of Meta’s own policies regarding use and collection of Facebook users’ data. Each of plaintiffs’ healthcare providers—MedStar Health System, Rush University System for Health, and UK Healthcare—allegedly installed the Meta Pixel on their patient portals. Plaintiffs claim that when they logged into their patient portal on their medical provider’s website, the Pixel transmitted certain information to Meta. They contend that this information, which is contemporaneously redirected to Meta, revealed their status as patients and was monetized by Meta for use in targeted advertising.]

care,²⁰ their deployment in healthcare must align with the major goals of improving care access, patient outcomes, and health equity.

Recommendation 1: Exercise oversight of federal agencies tasked with ensuring the responsible use of AI in the health care sector under the Biden AI Executive Order.

President Biden’s Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence²¹ (“AI EO”) sets out a comprehensive mandate to establish AI guardrails and establish the federal government as a model of accountable AI development and use.

The Executive Order states:

*Artificial Intelligence policies must be consistent with my Administration’s dedication to advancing equity and civil rights. My Administration cannot—and will not—tolerate the use of AI to disadvantage those who are already too often denied equal opportunity and justice. From hiring to housing to healthcare, we have seen what happens when AI use deepens discrimination and bias, rather than improving quality of life. Artificial Intelligence systems deployed irresponsibly have reproduced and intensified existing inequities, caused new types of harmful discrimination, and exacerbated online and physical harms.*²²

The Executive Order established Federal authority to “enforce existing consumer protection laws and principles and enact appropriate safeguards against fraud, unintended bias, discrimination, infringements on privacy, and other harms from AI. Such protections are especially important in critical fields like healthcare . . . where mistakes by or misuse of AI could harm patients, cost consumers or small businesses, or jeopardize safety or rights.”²³

Section 8 of the AI EO addresses the obligations of the federal government for “Protecting Consumers, Patients, Passengers, and Students.”²⁴ Notably, the AI EO encourages independent regulatory agencies, as they deem appropriate, to use their full range of authorities to protect American consumers from fraud, discrimination, and threats to privacy and to address other risks that may arise from the use of AI and to consider rulemaking, as well as emphasizing or clarifying where existing regulations and guidance apply to AI.²⁵ To achieve such purpose, the AI EO requires the relevant agencies to “clarify the responsibility of regulated entities to conduct due diligence on and monitor any third-party AI services they use, and emphasizing or clarifying requirements and expectations related to the transparency of AI models and regulated entities’ ability to explain their use of AI models.”²⁶

The AI EO directs the Secretary of the Health and Human Service (HHS) to establish an HHS AI Task Force, “to develop a strategic plan that includes policies and frameworks—possibly including regulatory action—on responsible deployment and use of AI and AI-enabled technologies in the health and human services sector (including research and discovery, drug and device safety, healthcare delivery and financing, and public health . . . to identify appropriate guidance and resources to promote that deployment including in the following areas, among others:

- Development, maintenance, and use of predictive and generative AI-enabled technologies in healthcare delivery and financing, considering appropriate human oversight of the application of AI-generated output; . . .
- Incorporation of equity principles in AI-enabled technologies used in the health and human services sector and monitoring algorithmic performance against discrimination and bias in existing models and helping to identify and mitigate discrimination and bias in current systems;
- Incorporation of safety, privacy, and security standards into the software-development lifecycle for protection of personally identifiable information, including measures to address AI-enhanced cybersecurity threats in the health and human services sector; . . .”²⁷

Chairman Wyden, you have stated, “As is frequently the case with new technology, AI provides us with exciting opportunities to better serve the American peo-

²⁰ HealthExec Report.

²¹ *Id.*

²² *Id.*, Section 2(d), pg. 75192.

²³ *Id.*, pg. 75192–3.

²⁴ *Id.*, pg. 75214.

²⁵ *Id.*

²⁶ *Id.*

²⁷ *Id.*

ple, but we're only beginning to see the consequences of leaving these systems unchecked. . . . The federal government has a responsibility to ensure the systems it is using to make decisions that impact Americans' daily lives are doing so accurately and without harmful bias."²⁸

We urge this Committee to exercise oversight on the actions to be completed by the HHS, specifically developing a robust strategy and identifying appropriate guidance and resources to ensure the responsible use of AI in healthcare and promote AI regulation based on the principles of non-discrimination, privacy protection, transparency, traceability, and contestability.

Recommendation 2: Move forward with comprehensive AI legislation

We commend the agency actions to guide the appropriate use of AI— including the April 2023 joint statement by the FTC, DOJ, EEOC, and CFPB on bias in automated systems,²⁹ and the White House Blueprint for an AI Bill of Rights.³⁰ However, since these measures alone do not address the full spectrum of governance required for the development, deployment, and use of AI systems, *we support the Algorithmic Accountability Act*³¹ and the *Blumenthal-Hawley Bi-Partisan Framework for U.S. AI Act*.³² We believe the implementation of either of these initiatives would go a long way in filling the legislative vacuum in which high-risk AI systems operate.

Specific to the use of AI in healthcare, we urge the Committee to consider the following recommendations in developing legislation:

- 1) AI systems should not be used for healthcare contexts where a certified clinical professional is required.
- 2) AI systems should not reinforce bias and discrimination by way of their application in healthcare and should respect data privacy rights of the citizens.
- 3) Healthcare data should be regulated to prevent the current loopholes regarding healthcare or wellness apps which do not fall under HIPAA.
- 4) AI-based health or wellness systems (*i.e.*, Fitbit) should not be used to make determinations for insurance or employment.
- 5) The consumer must be provided with clear, transparent, and complete information about AI-driven decision-making processes, including the subjects and elements involved. Simple processes should be in place for appealing any such decision.
- 6) Affected parties should have a right to contest adverse decisions made by AI systems.
- 7) Implement ex-ante impact assessments and ex-post evaluation or audit mechanisms for any AI system that implicates civil rights or public safety.

Given the serious challenges, we need federal legislation that mandates algorithmic transparency and accountability. We endorse the Hawley-Blumenthal bipartisan AI Act, a comprehensive framework for the governance of AI and urge this Committee to support the Algorithmic Accountability Act 2023 in moving forward to mark-up.

Thank you for your consideration of our views. We ask that this statement be included in the hearing record. We would be pleased to provide you and your staff with additional information.

Sincerely yours,

Marc Rotenberg
Executive Director

Merve Hickok
President

²⁸ U.S. Senate Committee on Finance, *Crapo, Wyden Press Federal Agencies on Use of Artificial Intelligence*, Newsroom (Nov. 9, 2023), <https://www.finance.senate.gov/ranking-members-news/crapo-wyden-press-federal-agencies-on-use-of-artificial-intelligence>.

²⁹ EEOC, CRT, FTC, and CFPB, *Joint Statement on Enforcement Efforts Against Discrimination and Bias in Automated Systems* (Apr. 25, 2023), https://www.ftc.gov/system/files/ftc_gov/pdf/EEOC-CRT-FTC-CFPB-AI-Joint-Statement%28final%29.pdf.

³⁰ White House Office of Science and Technology Policy, *Blueprint for an AI Bill of Rights* (October 2022), <https://www.whitehouse.gov/wp-content/uploads/2022/10/Blueprint-for-an-AI-Bill-of-Rights.pdf>.

³¹ Office of Senator Cory Booker, *Booker, Wyden, Clarke Introduce Bicameral Bill to Regulate Use of Artificial Intelligence to Make Critical Decisions Like Housing, Employment and Education*, Press Release (Sept. 21, 2023), <https://www.booker.senate.gov/news/press/booker-wyden-clarke-introduce-bicameral-bill-to-regulate-use-of-artificial-intelligence-to-make-critical-decisions-like-housing-employment-and-education>.

³² Senator Richard Blumenthal & Senator Josh Hawley, *Bipartisan Framework for U.S. AI Act*, <https://www.blumenthal.senate.gov/imo/media/doc/09072023bipartisanaiframework.pdf>.

Christabel Randolph
Law Fellow

Md Abdul Malek
Research Assistant

CONNECTED HEALTH INITIATIVE
1401 K St., NW, Ste. 501
Washington, DC 20005

February 16, 2024

The Honorable Ron Wyden
Chairman
Senate Finance Committee
Washington, DC 20515

The Honorable Mike Crapo
Ranking Member
Senate Finance Committee
Washington, DC 20515

RE: Statement for the Record of Brian Scarpelli, executive director of the Connected Health Initiative, on the hearing *Healthcare and AI: Promises and Pitfalls*, February 8, 2024

Dear Chairman Wyden, Ranking Member Crapo, and members of the Committee:

Thank you for holding this hearing on the topic of healthcare and artificial intelligence (AI). As you rightly note, there are both promises that we must realize and pitfalls to avoid when considering the use of AI systems in health contexts. Congress must ensure that federal healthcare policy enables patients and caregivers to leverage responsibly-designed AI to its full potential.

The American healthcare system desperately needs support. First and foremost, the American population is aging, and life expectancy is increasing, with those 65 or older accounting for one out of every five Americans by 2030—and 80 percent of those having at least one chronic condition. Second, healthcare costs are increasing, having already risen to more than approximately \$4.3 trillion annually, representing at least 17 percent of the U.S. gross domestic product. Finally, and no less troubling, the healthcare workforce is experiencing a growing shortage, with 30 out of the 35 physician specialties projected to experience serious deficits by the 2030s, with rural areas facing the brunt. The efficiencies AI offers are vital to overcoming these challenges.

The Connected Health Initiative (CHI) is a coalition of stakeholders dedicated to responsibly harnessing the power of technology to improve patient engagement and health outcomes. We advocate for policies that will improve patient outcomes, lower healthcare costs, improve the work-life balance of providers, and enable the continuing technological revolution across the healthcare ecosystem.

The quadruple aim

AI uses in healthcare range from back-office support and scheduling help to clinical decision support, encompassing a wide range of risk levels and types. We urge Congress and other policymakers to view the proven capabilities of AI in healthcare to assist patients, providers, technology developers, and others throughout the healthcare ecosystem through the lens of “quadruple aim” framework:

1. *Enhance Patients’ Outcomes and Experiences.* AI-supported interventions and treatments offer the ability to improve individual patient outcomes and engagement. Further, AI-supported engagement tools can help patients take steps to prevent disease and to stay engaged in their care after diagnosis.
2. *Improve Overall Population Health.* AI can help surface trends and suggest courses of action to address emerging or persistent health issues in population sets. An especially tricky aspect of population health management is that unforeseen factors—which are sometimes aspects of social determinants of health (SDOH) rather than traditional health indicators—often produce significant effects on the health of a given demographic. This is where large datasets and powerful analytical tools can equip providers and public health officials with actionable information and analyses.
3. *Reduce Costs.* AI capabilities have already been shown to be critical tools in maximizing efficiencies across the healthcare value chain. AI tools are starkly needed to assist in reducing administrative costs, as well as in leveraging data collected from outside of the doctor’s office between visits to support timely care plan updates and interventions that save the health system significant costs.
4. *Improve Healthcare Professionals’ Experience.* The healthcare sector is already experiencing a workforce shortage, with today’s frontline clinicians at high risk of burnout. AI-supported tools offer the opportunity to improve healthcare pro-

professionals' experience by maximizing efficiencies and capabilities, allowing them to reach more patients with better care. Yet, while AI tools are increasing job satisfaction and reducing burnout, they are not intended to replace the provider.

AI policy principles

As you, the federal agencies, states, and the private sector continue to explore how to responsibly bring AI to the healthcare system, CHI recommends that policy-makers align with CHI's cross-sectoral consensus health AI policy principles, which recognize the opportunities and challenges AI provides to healthcare and provide baseline recommendations across key areas including quality assurance/oversight, thoughtful design, access and affordability, and bias detection/mitigation, and data privacy and security, among other critical areas.¹ Other discrete opportunities for Congress and the federal government include:

- Leveraging consensus medical AI terminology² and CHI's cross-sectoral consensus understanding of the unique roles and interdependencies/shared responsibilities amongst the healthcare AI value chain³ as a baseline for the government's approach to health AI;
- Building on the leading efforts of the National Institute of Standards and Technology's voluntary AI Risk Management Framework⁴ and CHI's health-specific recommendations on its development and implementation⁵ to ensure that a coordinated approach is taken to health AI that scales risk mitigation requirements to intended uses and known harms;
- Helping build trust amongst providers and patients by enhancing transparency (supporting the sharing of information about the AI's intended use, development, performance, etc.) consistent with CHI's recommendations in *Advancing Transparency for Artificial Intelligence in the Healthcare Ecosystem*;⁶
- Advancing overdue Medicare coverage and payment policy changes that appropriately categorize AI (e.g., recognize that AI software as a medical device is appropriately categorized and paid for as a direct practice expense);
- Responsibly expanding support for its use in the prevention and treatment of beneficiaries' acute and chronic conditions;
 - Payment and coverage for health care AI systems must be informed by real world workflow and human-centered design principles; enable physicians to prepare for and transition to new care delivery models; support effective communication and engagement between patients, physicians, and the health care team; seamlessly integrate clinical, administrative, and population health management functions into workflow; and seek end-user feedback to support iterative product improvement.
- Ensuring that AI can support the transition to value-based care (e.g., eliminating barriers to the responsible use of health AI and other innovative technologies in the Merit-based Incentive Payment System and in Advanced Alternative Payment Models), consistent with recommendations in CHI's *Leveraging Digital Health to Realize Value-Based Care*.⁷

A risk-based approach

Policy frameworks should utilize risk-based approaches to ensure that the use of AI in healthcare aligns with recognized standards of safety, efficacy, and equity. Providers, technology developers and vendors, health systems, insurers, and other stakeholders all benefit from understanding the distribution of risk and liability in building, testing, and using healthcare AI tools. Policy frameworks addressing liability should ensure the appropriate distribution and mitigation of risk and liability. Specifically, those in the value chain with the ability to minimize risks based on their knowledge and ability to mitigate should have appropriate incentives to do so.

CHI has developed a framework for roles and interdependencies in AI regulation, which provides a guide for understanding the roles of different stakeholders in the

¹CHI Health AI Task Force's deliverables are accessible at <https://actonline.org/2019/02/06/why-does-healthcare-need-ai-connected-health-initiative-aims-to-answer-why/>.

²E.g., <https://www.ama-assn.org/practice-management/cpt/cpt-appendix-s-ai-taxonomy-medical-services-procedures>.

³<https://connectedhi.com/wp-content/uploads/2024/02/CHI-Health-AI-Roles.pdf>.

⁴<https://www.nist.gov/itl/ai-risk-management-framework>.

⁵<https://actonline.org/wp-content/uploads/Policy-Principles-for-AI.pdf>.

⁶CHI's recommendations on necessary policy changes to enhance transparency for healthcare AI are available at <https://bit.ly/3Gd6cxs>.

⁷<https://connectedhi.com/wp-content/uploads/2022/02/LeveragingDigitalHealth.pdf>.

AI value chain. Regulation like this would help foster a shared sense of responsibility for AI development. The framework is appended to this testimony.

Leveraging existing authority

Many of the concerns that advocates have with AI regulation can be addressed by existing authority within various agencies. The Food and Drug Administration (FDA) already has statutory authority to assess medical devices for safety, including devices that use AI. Other agencies, including the Federal Communications Commission (FCC) and Federal Trade Commission (FTC) apply their existing authority to AI scenarios. CHI encourages agencies to apply their statutory authority to AI contexts as Congress continues to develop a more nuanced framework for handling AI policy.

The WEAR IT Act

To further the adoption of digital medicine and responsible AI in healthcare, CHI supports H.R. 6279, the Wearable Equipment Adoption, Reinforcement, and Investment in Technology (WEAR IT) Act. The WEAR IT Act would allow individuals to access certain wearable health technology through their tax-advantaged flexible spending accounts (FSAs) and health savings accounts (HSAs). Currently the IRS allows HSA and FSA funds to be spent primarily on single-purpose devices. In a recent development, the IRS now considers the Oura Ring and the Aura Pulse Comprehensive Health Tracker eligible for FSA and HSA expenditures, two exceptions to the IRS's general rule against such devices. Many cutting-edge wearable health devices have multiple functions such as catastrophic fall detection, heart rate monitoring, and/or blood oxygen measuring. Although these devices outperform covered legacy technology in many cases, they are generally not covered (with the exceptions described above) because of the IRS's historical interpretation of the law, which is outdated. The IRS has recently begun to modernize its approach to HSA and FSA eligibility. However, if the WEAR IT Act is enacted, such devices will be covered by FSAs and HSAs, giving consumers more choice and additional ways to pursue their health goals. Moreover, the integration of AI into these devices can also prove to be helpful to consumers as it can help to simplify function and provide more detailed information for both the users and healthcare providers.

We are seeing the positive effects that AI can have in healthcare and as AI technology advances, it is important to promote—and eliminate barriers to—its responsible use while ensuring AI is safe and effective. We need to ensure that patients are getting the care that they need while also making sure that healthcare workers are not overburdened to provide high-quality care. We ask the Committee to take our points into consideration as policymakers seek avenues to ensure AI's continued, responsible integration into healthcare.

Sincerely,

Brian Scarpelli
Executive Director

AI Roles and Interdependencies Framework

CHI urges all stakeholders in the healthcare ecosystem that are developing and using AI to align with CHI's consensus health AI principles, which recognize the shared responsibility for AI safety, efficacy, and transparency. CHI supports (1) leveraging a risk-based approach to AI harm mitigation where the level of review, assurance, and oversight is proportionate to potential harms and (2) those in the value chain with the ability to minimize risks based on their knowledge and ability, and having appropriate responsibilities and incentives to do so. Further, managing AI/Machine Learning (ML) risks will be more challenging for small to medium-sized organizations, depending on their capabilities and resources. Building on these general health AI principles, CHI proposes clear definitions of stakeholders across the healthcare AI value chain, from development to distribution, deployment, and end use. Then, CHI suggests roles for supporting safety, ethical use, and fairness for each of these important stakeholder groups that are intended to illuminate the interdependencies between these actors, thus advancing the shared responsibility concept. These roles and interdependencies are also mapped to the Functions defined in the National Institute of Standards and Technology's (NIST's) AI Risk Management Framework (RMF).

Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
AI/ML Developers	Someone who designs, codes, re-searches, or produces an AI/ML system or platform for internal use or for use by a third party. See below for defined Subgroups of this Stakeholder Group along with recommendations specific to that Subgroup.	- Informing deployers and users of data requirements/definitions, intended use cases/populations and applications (<i>e.g.</i>, disclosing sufficient detail allowing providers to determine when an AI-enabled tool should reasonably apply to the individual they are treating), including whether the AI/ML tools are intended to augment human work versus automate workflows, and status of compliance with all applicable legal and regulatory requirements. - Prioritizing safety, efficacy/accuracy, transparency, data privacy and security, and equity from the earliest stages of design, leveraging (and, where appropriate, updating) existing medical AI/ML guidelines on research and ethics, leading standards, and other resources as appropriate. - Employing algorithms that produce repeatable results and, when feasible, are auditable, and make decisions that (when applied to medical care) are clinically validated, fostering efficacy through continuous monitoring. - Utilizing risk management approaches that scale to the potential likely harms posed in intended use scenarios to support safety, protect privacy and security, avoid harmful outcomes due to bias, etc. - Providing information that enables those further down the value chain can assess the quality, performance, equity, and utility of AI/ML tools. - Aligning with relevant ethical obligations and international conventions on human rights and supporting the development of new ethical guidelines to address emerging issues as needed.	AI Deployment; Operation and Monitoring; Test, Evaluation, Verification, and Validation (TEVV); Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight

Stakeholder Subgroup	Definition	Roles
Foundation Model Developer	Someone who creates or modifies large and generalizable machine learning models that can be used/adapted for various downstream tasks and applications, such as natural language processing, computer vision, or software development.	Building on the cross-AI/ML Developer roles noted above: - Assessing what bias and safety issues might be present in its Foundation Model, and documenting steps taken to mitigate those issues in its Transparency Documentation (e.g., Transparency Notes, System Cards and product documentation). - Providing clear guidance on (1) how to use and adapt its Foundation Model for various foreseeable downstream tasks and applications, and (2) what limitations or risks may arise from doing so based on challenges discovered during testing and deployment.
AI Platform Developer	Someone who leverages existing foundation models and builds an industry-agnostic platform that enables other developers to access, customize, and deploy these models for various use cases and applications, such as natural language processing, computer vision, and/or software development.	Building on the cross-AI/ML Developer roles noted above: Testing for, identifying, and mitigating bias and safety issues that may arise from using or modifying existing foundation models for its AI Platform, and documenting these issues and steps taken to address them in its transparency documentation (e.g., transparency notes, system cards and product documentation).
Health AI Platform Developer	Someone who creates or uses AI-powered platforms that are tailored for the healthcare domain, such as administrative efficiency, diagnostics, therapeutics, or research. These platforms may leverage foundation models (or other types of machine learning models or solutions), such as AI platforms, that are suitable for specific healthcare problems and data sources.	Building on the cross-AI/ML Developer roles noted above: - Meeting specific requirements and standards of the healthcare domain, such as accuracy, efficacy, explainability, and compliance with regulations. - Testing for, identifying, and mitigating any bias and safety issues that may affect the health outcomes of patients or the performance of clinicians using the Health AI Platform, and documenting these issues and the steps it has taken to address them in its transparency documentation (e.g., transparency notes, system cards and product documentation).

Stakeholder Subgroup	Definition	Roles
Digital Health Solution Developer	Someone who creates complete digital tools and technologies to improve health and healthcare outcomes, such as providing diagnostic and administrative solutions for clinicians, patients, and healthcare organizations. They may build digital health solutions with both health AI platforms, which are specialized for the health care domain, and AI platforms, which are more general and adaptable for various use cases and applications.	Building on the cross-AI/ML Developer responsibilities noted above: • Specifying appropriate uses for its digital health solution to avoid amplifying bias or safety issues that may exist in the underlying foundation models, AI platforms, or health AI platforms. • Designing user interfaces to enable an end user to safely and effectively act upon the output of the tool, such as providing explanations, feedback mechanisms, or human oversight options, providing clear documentation to Deploying Organizations and Users to help them avoid bias and safety issues.

Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
Deploying Organization (Healthcare Provider or Payor)	Someone who is a healthcare providers and health care payors that and is deploying solutions built by Digital Health Solution Developers. They may also have their own internal IT staff that use health AI platforms or general AI platforms to develop their own custom digital health solutions.	<i>Respecting that managing AI/ML risks will be more challenging for small to medium-sized organizations depending on their capabilities and resources:</i> • Adopting AI/ML Developer instructions for use, specifying appropriate uses for Users through governance policies to avoid bias and safety issues that may exist in the underlying foundation models, AI platforms, or health AI platforms. • Developing and leveraging digital health solutions that augment efficiencies in coverage and payment automation, facilitate administrative simplification/reduce workflow burdens, and are fit for purpose. • Setting organization policy/designing workflows to reduce the likelihood that a User will act upon the output of the tool in a way that would cause fairness/bias or safety issues (tailored explanations, feedback mechanisms, and/or human oversight options). • Developing and organizational guidance on how the digital health solution should and should not be used. • Creating risk-based, tailored communications and engagement plans to enable easily understood explains to patients about how the digital health solution was developed, its performance and maintenance, and how it aligns with the latest best practices and regulatory requirements.	AI Deployment; Operation and Monitoring; Domain Expert; AI Impact Assessment; Procurement; Governance and Oversight

Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
Provider/Clinician Users and Administrative Users	Someone who directly interacts with or benefits from the digital health solutions that are built by Digital Health Solution Developers or by the internal IT staff of the Deploying Organization. They may include clinicians, such as doctors, nurses, or pharmacists, and administrative staff, such as billing, claims, or customer service personnel, in the provider and payor organizations.	<i>Respecting that managing AI/ML risks will be more challenging for small to medium-sized organizations depending on their capabilities and resources:</i> • Taking required training and incorporating employer guidance about use of AI/ML digital health solutions. • Documenting (through automated processes or otherwise) whether AI is being used in medical records and report any issues or feedback to the developer, such as errors, vulnerabilities, biases, or harms (where AI/ML's use is known by the User). • Ensuring there is appropriate clinician review and review of the output or recommendations from each digital health solution prior to acting on it (where AI/ML's use is known by the User).	AI Deployment; Operation and Monitoring; Domain Expert; AI Impact Assessment; Procurement; Governance and Oversight

Payer Users (Centers for Medicare and Medicaid Services [CMS], State Medicaid, Private)	Someone that pays for the cost of healthcare services administered by a healthcare provider.	• Leveraging AI/ML systems that improve efficiencies in coverage and payment automation, facilitate administrative simplification, and reduce provider workflow burdens. • Aligning with medical AI/ML definitions, present-day and future AI/ML solutions, the future of AI/ML medical coding changes and trends. • Developing support mechanisms for the use of AI/ML by providers based on clinical validation, aligning with clinical decision-making processes familiar to providers, and high-quality clinical evidence. • Assuring that AI/ML systems allow for the individualized assessment of specific medical and social circumstances and provider flexibility to override automated decisions, ensuring that use of AI/ML does not improperly reduce or withhold care, or overrides the provider's clinical judgement. • Disclosing information about training and reference data to demonstrate that AI/ML systems do not create or exacerbate inequities and that protections are in place to mitigate bias. • Developing and proliferating easy to understand resources for beneficiaries and their providers that capture how and when AI/ML is being used, what information it is leveraging, and what it means to patients.	AI Deployment; Operation and Monitoring; Domain Expert; AI Impact Assessment; Procurement; Governance and Oversight
Patient Groups/ Patient Users	Someone who uses digital tools and technologies that are built by Digital Health Solution Developers or experiences their use in treatment.	• Developing and proliferating easy to understand resources that capture how AI/ML is being used and what it means to patients/patient groups, including explanations on the purpose and limitations of the digital health solutions that they use or benefit from (e.g., diagnostic, therapeutic, administrative). • Raising awareness of patients' rights and choices when using digital health solutions, such as consent, access, correction, or deletion of their personal data.	Human Factors

Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
Standard-Setting Organizations	An organization whose primary function is developing, coordinating, promulgating, revising, amending, reissuing, interpreting, or otherwise contributing to the usefulness of technical standards to those who employ them.	Developing and promoting adoption of international voluntary/non-regulatory consensus standardized approaches and resources to steward a shared responsibility approach to AI.	Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight
Certification Bodies & Test Beds	A certification body is a third-party organization that assures the conformity of a product, process or service to specified requirements. A test bed is a platform for conducting rigorous, transparent, and replicable testing of scientific theories, computing tools, and new technologies to a standard.	- Creating and making available transparent and reliable processes for the assurance of conformity to voluntary AI standards. - Creating and making available voluntary sandbox environments to help evaluate the usability and performance of AI/ML-based high-performance computing applications to advance the understanding of how reliable and efficacious AI, and to provide an appropriate assurance of reliability and efficacy.	Test, Evaluation, Verification, and Validation (TEVV); Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight
Accrediting and Licensing Bodies, and Medical Specialty Societies and Boards	Accrediting and licensing bodies are governing authorities that establish the suitability of any participating certification body. Notably, state-level board serve this purpose for physicians, nurses, and other clinicians to standards set by each state. Medical specialty societies are organizations for physicians, research and clinical scientists who are actively involved in the study of a particular specialty.	- Based on clinical needs and expertise, developing and setting the medical standard of care and ethical guidelines to address emerging issues with the use of AI/ML in healthcare needed to advance the quadruple aim. - Identifying the most appropriate uses of AI-enabled technologies and developing and disseminating guidance and education on the responsible deployment of AI/ML in healthcare, both generally and for specialty-specific uses.	Test, Evaluation, Verification, and Validation (TEVV); Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight

Academic and Medical Education Institutions	Tertiary educational institutions, professional schools, or forms a part of such institutions, that teach medicine and awards a professional degree for physicians or other clinicians.	• Developing and teaching curriculum that will advance understanding of and ability to use healthcare AI/ML solutions responsibly, which should be assisted by inclusion of non-clinicians such as data scientists and engineers as instructors. • Developing curriculum to advance the understanding of data science research to help inform ethical bodies (<i>e.g.</i>, Institutional Review Boards that are reviewing protocols of clinical trials of AI/ML-enabled medical devices).	Human Factors; Domain Expert; AI Impact Assessment
--	--	--	---

FEDERATION OF AMERICAN HOSPITALS

750 9th Street, NW, Suite 600
 Washington, DC 20001
 202-624-1500
 FAX 202-737-6462
<https://www.fah.org/>

United States Senate
 Committee on Finance

Re: “Artificial Intelligence and Health Care: Promises and Pitfalls”

The Federation of American Hospitals (FAH) submits the following Statement for the Record in response to the Senate Finance Committee’s (Committee’s) full committee hearing “Artificial Intelligence: Promises and Pitfalls.” Hospitals are on the front lines of utilizing technology and innovation to improve patient experience, care, and outcomes. We appreciate the Committee’s efforts to understand the use of algorithms and artificial intelligence (AI) systems in health care.

The FAH is the national representative of more than 1,000 tax-paying community hospitals and health systems throughout the United States. FAH members provide patients and communities with access to high-quality, affordable care in both urban and rural areas across 46 states, plus Washington, DC, and Puerto Rico. Our members include teaching, acute, inpatient rehabilitation, behavioral health, and long-term care hospitals and provide a wide range of inpatient, ambulatory, post-acute, emergency, children’s, and cancer services.

Hospitals are consistently at the forefront of innovation and technology. Our members have seen firsthand the potential AI has to unlock efficiency and improve delivery through management and orchestration of patient care. Hospitals utilize AI for a wide range of activities including enhancing clinical documentation, streamlining nurse handoff processes, and optimizing staffing, scheduling, and improving care delivery.

Care delivery, in particular acute care, is a complex matrix of activities that requires coordination and engagement between numerous stakeholders (*e.g.*, nurses, physicians, pharmacists, technicians, patients, and family members). AI is capable of understanding this complexity and orchestrating care delivery by identifying the next best action to take at any step in the process, ensuring precious time and resources are deployed in the most efficient and effective manner. For example, regarding bed management, AI could unlock bottlenecks in a hospital through understanding the interdependencies of moving patients within a hospital (*e.g.*, from the emergency room to an inpatient unit). It could select the correct next action (which patient moves next and where) that balances the needs of the patients, bandwidth and proficiency of the staff, and geography of the care teams assuming responsibility. Optimizing these decisions can unlock significant capacity in hospitals on a daily basis and improve patient experience.

Policymakers should encourage the use of generative AI specifically designed to simplify access, consumption, and readability of health care data. For example, a voice assistant for clinicians to extract specific information from large patient history can encourage evidence-based decisions and reduce clinical errors. Generative AI tools can contextually summarize, sequence, and modularize data better than static digital systems.

We urge Congress, regulators, developers of AI tools, and users of AI tools to collaborate on appropriate frameworks to maximize the safety and efficacy of AI within the health care sector. We note that a layered approach would be most appropriate, and legislative and regulatory policy should distinguish between whether a health care system or organization is developing their own platform and tools for internal use, creating a commercial product, or contracting with an outside vendor for internal use of a commercial product. It also should differentiate between AI uses to augment human decision making versus a situation where the output of algorithms is patient facing or directing patient care.

Further, guardrails should address topics such as transparency, ethical use, and oversight. The developers of commercial products that embed AI tools should make measures available that address issues such as how a model works, the data used to train it, appropriate and inappropriate uses of the tool, and results of any testing that have been done to assess bias to ensure that AI’s use in care models decreases disparities and does not exacerbate them. Guardrails also should address ethical use, *i.e.*, when it is necessary to have “a human in the loop,” as well as the oversight of AI tools that are in use to ensure that they are functioning appropriately. The

most effective way to regulate AI is to focus on the processes by which AI is developed, rather than on the individual algorithms themselves. When companies and organizations are developing AI products using a trusted process and framework, they have more ability to innovate new products and versions while mitigating risk. By focusing on the processes by which AI is developed, we can ensure that AI is developed in a manner that is safe and ethical.

Finally, Congress also should give thoughtful consideration to the topic of liability, which is a new and challenging aspect of AI. While health care providers bear responsibility for the care they provide, the developers of commercial AI products must also be accountable if safety, bias, or other harms are caused by a flaw in the AI tool itself.

We look forward to working with the Committee on these critical issues. If you have any questions or would like to discuss these comments further, please do not hesitate to contact me or a member of my staff at (202) 624-1534.

Sincerely,

Charles N. Kahn III
President and CEO

HEALTHCARE CONFIDENTIALITY COALITION

February 5, 2024

Senator Ron Wyden
Chairman
U.S. Senate
Committee on Finance
219 Dirksen Senate Office Building
Washington DC 20510-6300

RE: Full Committee Hearing on “Artificial Intelligence and Health Care: Promise and Pitfalls,” February 8, 2024

Dear Chairman Wyden:

The Healthcare Confidentiality Coalition (Coalition) thanks you and other members of the U.S. Senate Committee on Finance (Committee) for holding a hearing titled “Artificial Intelligence and Health Care: Promise and Pitfalls” on February 8, 2024. We appreciate the opportunity to submit this statement for the record.

The Healthcare Confidentiality Coalition is a diverse group of health organizations committed to advancing effective and workable health information privacy, security, and interoperability policies at the federal and state levels. Our mission is to advocate for policies and practices that safeguard the privacy and security of patients and healthcare consumers while, at the same time, enabling the essential flow of patient information that is critical to the timely and effective delivery of healthcare, improvements in quality and safety, and the development of new lifesaving and life-enhancing medical interventions. Our members include hospitals, health systems, medical teaching colleges, health plans, pharmaceutical companies, medical device manufacturers, vendors of electronic health records, biotech firms, employers, health product distributors, pharmacies, pharmacy benefit managers, health information and research organizations, and others. Through the diversity of our membership, we are able to develop a nuanced perspective on the impact of legislation, regulation, and other policies affecting the critical flow of essential health information.

The hearing could not be timelier or more aptly titled. We strongly believe that artificial intelligence (AI) is a transformational tool that holds enormous promise for health care, but at the same time presents serious potential pitfalls and risks. Guardrails are essential to protect against these risks, and we support the efforts by the Committee and others in Congress to establish a regulatory framework that will appropriately address these risks. This regulatory framework must be carefully calibrated to strike the right balance between protecting individual rights and patient safety and encouraging competition and innovation.

As outlined in the Coalition’s Principles for the Responsible Development and Use of Artificial Intelligence in Health Care (AI Principles), we support an approach to AI regulation that considers a national data privacy standard as a foundational element to ensure the responsible and beneficial use of AI. We have long been on record calling for national privacy legislation to protect personal health data, as re-

flected in our Beyond HIPAA Principles. While the scale and sophistication of the collection, use and generation of data as part of technological development in all areas has made the need for a national data privacy standard an imperative for the advancement of society, this is particularly so in AI, which relies on massive amounts of data to power its learning.

It is only when individuals have the assurance and confidence that their personal data will be use appropriately and responsibly in accordance with their reasonable expectations, and safeguarded against misuse or misappropriation, that they will embrace technologies such as AI, that depend so heavily on this data. Thus, the principles of privacy by design should be integrated into AI tools from the start. This includes, but should not be limited to, data minimization, use limitations and individual rights, such as the right to know, access, correct and, if feasible, delete personal information.

Similarly, security safeguards, which may be based on guidelines such as those provided in the National Institute of Standards and Technology (NIST) Cybersecurity Framework and Risk Management Framework, must be implemented to protect against data breaches and other threats that could expose the data used or alter the use, behavior, or performance of an AI application. Any regulatory framework should take a risk-based approach that considers potential impact and possible harms. By focusing on a risk-based approach and regulating accordingly, regulators will allow developers and users to allocate resources appropriately and proportionate to the potential harm and nuances of their specific AI use case.

Thank you for your consideration of our comments. Please do not hesitate to contact me at tgrande@confcoa.org or 202-750-1989 if you have any questions.

Sincerely,

Tina O. Grande
President and CEO

HEALTHCARE LEADERSHIP COUNCIL

February 8, 2024

The Honorable Ron Wyden
Chairman
Senate Finance Committee
Dirksen Senate Office Building
Washington, DC 20510

The Honorable Mike Crapo
Ranking Member
Senate Finance Committee
Dirksen Senate Office Building
Washington DC 20510

Dear Chairman Wyden and Ranking Member Crapo:

The Healthcare Leadership Council (HLC) applauds your efforts to ensure innovation and safety in the development and use of artificial intelligence for healthcare.

HLC is a coalition of chief executives from all disciplines within American healthcare. It is the exclusive forum for the nation's healthcare leaders to jointly develop policies, plans, and programs to achieve their vision of a 21st century healthcare system that makes affordable high-quality care accessible to all Americans. Members of HLC—hospitals, academic health centers, health plans, pharmaceutical companies, medical device manufacturers, laboratories, biotech firms, health product distributors, post-acute care providers, homecare providers, group purchasing organizations, and information technology companies—advocate for measures to increase the quality and efficiency of healthcare through a patient-centered approach.

Our members are already familiar with the positive effect AI can have in healthcare. For decades traditional rule-based AI has been a tool not only to aid administrative burden by automating processes like medical billing and recommending codes, but also to provide clinical decision support and establish clinical care standards to improve quality of care.¹ For example, the use of AI in the review and translation of mammograms is 30 times faster than just a physician with 99% accuracy, reducing the need for patients to undergo unnecessary biopsies and diminishing waste within the sector.²

¹Exploratory Study of Artificial Intelligence in Healthcare (January 2016), Exploratory Study of Artificial Intelligence in Healthcare-libre.pdf (dlwqtxts1xzle7.cloudfront.net).

²See <https://www.pwc.com/gx/en/industries/healthcare/publications/ai-robotics-new-health/transforming-healthcare.html>, referencing the California Biomedical Research Association. New

Now with the emergence of generative AI capable of identifying patterns and improving upon user identified if-then procedures, our members see an even greater potential for the technology to shape advancements across the entire healthcare sector. Researchers are able to integrate AI into the fabric of drug discovery to identify leads faster than ever before. Pharmacies are able to utilize market data to anticipate demand so commonly used medications remain in stock at patients' preferred pharmacies, strengthening their role in the supply chain. Clinicians are even able to automatically distill their conversations with patients into easy-to-understand notes and instructions, ensuring the increase in care quality is not lost in translation.

The greatest obstacle towards realizing the potential of AI in healthcare is building trust in its implementation. As with other transformative technologies, these tools will need to be reasonably regulated, and guardrails are essential to protect patients from risks that emerge when working with health and personal data. It is our hope the federal government collaborate with all stakeholders to strike the right balance between protecting individual liberties and allowing this continued innovation to flourish so the full promise of AI in healthcare can be utilized. In collaboration with our members and the Healthcare Confidentiality Coalition we have developed a series of principles to guide future regulatory framework.

Addressing Adverse Bias and Discrimination

AI applications in health care present the risk of bias as the underlying, especially historical, data sets may lack representative or accurate data. Access to high quality data sets that are as complete as possible, including sensitive personal information (e.g., data on race, ethnicity, gender, etc.), is ideal, but not always possible. Organizations should then take comprehensive steps to identify and mitigate potential sources of harmful bias across the lifecycle of their model development, and where reasonable and appropriate for specific models, align with industry-developed standards. It is important that this not be done by excluding sensitive personal data or data of vulnerable groups from AI training data so that any bias may be detected and remedied, and all patients may benefit from the advances in health care brought about by AI.

Risk-based Approach

Regulatory agencies and organizations using AI applications should take a risk-based approach to the regulation and oversight of AI applications, taking into account their potential impact and possible harms. Organizations should perform risk assessments that align with, or extend beyond, consensus-based risk management frameworks such as the AI Risk Management Framework (AI RMF) developed by the NIST. An AI Risk Assessment should identify potential risks that the AI tool could introduce, potential mitigation strategies, detailed explanations of recommended uses for the tool, and risks that could arise should the tool be used inappropriately.

By focusing on a risk-based approach and regulating accordingly, regulators will allow developers and users to allocate resources appropriately and proportionate to the potential harm and nuances of their specific AI use case. Healthcare applications with the highest risk, for example, would in turn have the highest guardrails, such as requiring more frequent review or human intervention. Additionally, regulators should avoid imposing duplicative compliance requirements, and consideration should be given to organizations that follow a framework such as the NIST AI RMF in the imposition of penalties.

Federal Standards

Any regulatory framework(s) for AI applications should be developed and applied at the federal level. A single national standard that preempts state laws in this area will avoid conflicting requirements and facilitate compliance without unduly restricting innovation.

Privacy and Security

Personal information used in AI should be subject to robust privacy and security protections at the federal level. This includes adhering to existing health data privacy and security protections in the Health Insurance Portability and Accountability Act of 1996 (HIPAA) for protected health information and equivalent protections for non-HIPAA health data. The principles of privacy by design should be integrated into AI tools from the start. This includes, but should not be limited to, data mini-

mization and use limitations. Individuals should have the right to be informed about the collection and use of their personal information, and the right to access, correct and, if feasible, delete their personal information. Congress should establish a single national standard for the use of personal information not already subject to HIPAA that includes standards for the use of that information in AI applications by entities not regulated by HIPAA. Security safeguards, which may be based on guidelines such as those provided in the National Institute of Standards and Technology (NIST) Cybersecurity Framework and Risk Management Framework, should protect against data breaches, data poisoning, exfiltration of models or training data and other threats that could expose the data used or alter the use, behavior, or performance of an AI application.

Harmonization

Federal agencies such as the Office for Civil Rights (OCR), Food and Drug Administration (FDA), the U.S. Equal Employment Opportunity Commission and the Federal Trade Commission (FTC), among others, should collaborate to align the federal government's approach to the regulation of AI. OCR and the FDA have worked together in the past to align (*e.g.*, on the regulation of medical devices) as do OCR and the FTC on health information privacy. This will allow organizations subject to the authority of different federal agencies to align their approach to implementing AI applications across the enterprise, avoiding confusion, and leading to greater compliance.

While each agency may approach the technology from a different regulatory angle, whether safety, privacy, consumer protection or otherwise, they should be able to take a patient-centered approach to reach sufficient alignment so that compliance with one framework will not result in violation of, or inability to comply with, another. Failure to harmonize regulatory frameworks will not only create interpretation and compliance burdens but will slow AI development and stifle innovation by creating a regulatory patchwork that fails to account for how health care is delivered. Additionally, conflicts between federal agencies' regulation of AI will hamper U.S. efforts to lead globally in the regulation of AI. Other countries want to adopt a framework for the regulation of AI that harmonizes across business sectors and regulatory areas, rather than having to deal with discordant or conflicting requirements.

Accountability

Health organizations that use AI should determine and establish a risk-based structure of accountability that extends across its partnerships to ensure that their AI use cases are deployed in a responsible, fair, and consistent manner. This includes developing, implementing, and documenting principles, policies, procedures, as well as an internal collaborative governance structure and controls to oversee the development and use of AI applications. These controls should include quality control parameters for the data used as well as criteria against which the performance of the AI applications is monitored, evaluated, and re-evaluated, as needed, at regular intervals throughout the lifecycle. Accountability should extend to the highest levels of management and should include key elements such as risk-assessment, training, monitoring and internal sanctions.

Transparency

Transparency is essential to build trust in AI technology. Where appropriate, organizations should disclose when they are using AI tools, especially when these tools are used to make decisions about individuals. Organizations should not be required to reveal the inner workings of their AI systems to the public or regulatory agencies, nor is there any benefit in doing so. The detailed disclosure of either data inputs or algorithmic processes would not be meaningful to patients, providers, or payers, would force AI developers to disclose their intellectual property or proprietary technology, could create AI vulnerability risks, and may limit innovators willingness to work with the already highly regulated healthcare industry on meaningful AI applications.

Explainability

Developers of AI applications for use in health care must be able explain to users how a decision is made by a high-impact AI application in a way that is sufficiently understandable. All stakeholders should be able to gauge the context in which an algorithm operates and understand the implications of the outcomes. Users should in turn be able to explain the role of algorithms to individuals affected by AI-assisted decisions. Explanations should be meaningful and useful, tailored to the audience and calibrated to the level of risk.

Thank you for your attention to this important matter. Should you have any questions, please contact Sam Carley at 202-449-3445 or scarley@hlc.org.

Sincerely,

Maria Ghazal
President and CEO

MEDICAL GROUP MANAGEMENT ASSOCIATION
1717 Pennsylvania Ave., NW, #600
Washington, DC 20006
T 202-293-3450
F 202-293-2787
<https://www.mgma.com/>

February 7, 2024

The Honorable Ron Wyden
Chairman
Senate Committee on Finance
219 Dirksen Senate Office Building
Washington, DC 20510

The Honorable Mike Crapo
Ranking Member
Senate Committee on Finance
219 Dirksen Senate Office Building
Washington, DC 20510

Re: MGMA Statement for the Record—Senate Committee on Finance’s February 8th hearing, “Artificial Intelligence and Health Care: Promise and Pitfalls”

Dear Chairman Wyden and Ranking Member Crapo:

On behalf of our member medical group practices, the Medical Group Management Association (MGMA) would like to thank the Committee for the opportunity to provide feedback in response to the February 8th hearing, “Artificial Intelligence and Health Care: Promise and Pitfalls.” We appreciate your leadership in examining the impact of artificial intelligence (AI) on the healthcare sector.

With a membership of more than 60,000 medical practice administrators, executives, and leaders, MGMA represents more than 15,000 group medical practices ranging from small private medical practices to large national health systems, representing more than 350,000 physicians. MGMA’s diverse membership uniquely situates us to offer the following policy recommendations.

Please find attached MGMA’s issue brief reviewing AI and our advocacy priorities for your consideration. We are happy to discuss these priorities at any time and look forward to collaborating with the Committee to craft sensible policies that adequately balance the promise of AI technology with its potential risks. If you have any questions, please contact James Haynes, Associate Director of Government Affairs, at jhaynes@mgma.org or 202-293-3450.

Sincerely,

Anders Gilberg
Senior Vice President, Government Affairs

MGMA 2024 ARTIFICIAL INTELLIGENCE ISSUE BRIEF

While certain Artificial Intelligence (AI) capabilities have long been around in the healthcare space, in recent years there has been a significant acceleration in the introduction and adoption of new AI technologies. This has led to increased congressional and regulatory consideration of how AI operates within the industry and how best to regulate its use. MGMA advocates for policies that bolster the development and utilization of effective and ethical AI tools to improve operational efficiencies for medical groups and support high-quality patient care.

AI is generally characterized as technology capable of simulating human thought and performing real-world tasks. Different organizations and government bodies use separate definitions that are context-specific but colloquially referred to as AI.

Predictive AI may use algorithms to analyze large amounts of data to make predictions, while generative AI is trained on large datasets to create new content. Machine learning technology can analyze large datasets for patterns and gain insights that are applied to decision making; natural language processing allows computers to understand and manipulate human language. All told, AI technology can take many forms.

USE OF AI IN HEALTHCARE

AI technology is utilized by medical groups in numerous facets of healthcare. AI-enabled tools can do everything from helping revenue cycle management by improving medical coding to providing predicative analyses of performance areas and assisting in patient communications and marketing efforts. New technologies have the potential to augment clinical decision making, as well as streamline operations and lower administrative costs.

Unfortunately, while AI affords much opportunity for positive change, there have been noteworthy examples of the technology being used to determinantal effect. Medical group practices have raised concerns that certain AI tools may be used to aggravate administrative burdens such as mass, rapid denials of prior authorization requests, large language models providing “hallucinations” or inaccurate answers, and more. AI could offer significant benefits to medical groups, but it’s important to understand the risks and have safeguards in place ahead of more widespread adoption.

RECENT ADMINISTRATIVE DEVELOPMENTS

The Biden Administration recently issued an Executive Order¹ focusing on the use of AI in many different sectors of the economy, including healthcare. There is no unified national policy to oversee the development and deployment of AI; various agencies are tasked with regulating certain sections of AI in healthcare such as the Food and Drug Administration (FDA) overseeing medical devices utilizing AI. The Executive Order signaled a coordinated approach from the federal government regarding instituting AI safeguards and included numerous directives to federal agencies.

In response to the Executive Order, 28 healthcare systems and payers—such as Geisinger, Mass General Brigham, and Sanford Health—signed a pledge² committing to align the industry around principles that AI should promote healthcare outcomes that are “Fair, Appropriate, Valid, Effective, and Safe (FAVES).” At least 15 leading companies that develop AI technology have since offered similar commitments to the White House.

The Office of the National Coordinator for Health Information Technology finalized a rule meant to increase AI transparency near the end of 2023. Specifically, the final rule established transparency requirements for AI and certain predictive algorithms that are part of certified information technology (IT). The agency’s approach was to ensure that users of certified health IT can access information about AI and predictive algorithms, and that the technology follows the FAVES principles. Other federal agencies have signaled their intent to issue federal regulations on AI in the near future.

CONGRESSIONAL ATTENTION

In an effort to better understand the technology, Congress has held numerous forums on AI such as closed-door briefings and hearings in anticipation of potentially introducing legislation to regulate the industry. Prominent executives from AI companies have testified about the potential oversight, while healthcare leaders have addressed both chambers on the benefits and challenges associated with AI programs.

Senate Committee on Health, Education, Labor, and Pensions (HELP) Ranking Member Bill Cassidy issued a white paper on AI’s use in healthcare and called for the public’s feedback on the regulation and development of AI. The white paper³ reviewed policy areas that could require updated laws and rules, while at the same time examining the possibility of AI to help develop new medicines, reduce the workload of healthcare providers, and more. This offers an indication of where congressional leaders are heading in terms of legislation.

ADVOCACY PRIORITIES

- **Medical groups should be able to easily understand the use and function of AI products;** whether they be a stand-alone product or integrated into

¹<https://www.whitehouse.gov/briefing-room/statements-releases/2023/10/30/fact-sheet-president-biden-issues-executive-order-on-safe-secure-and-trustworthy-artificial-intelligence/>.

²<https://www.whitehouse.gov/briefing-room/blog/2023/12/14/delivering-on-the-promise-of-ai-to-improve-health-outcomes/>.

³https://www.help.senate.gov/imo/media/doc/help_committee_gop_final_ai_white_paper1.pdf.

- other technology. Proper transparency and disclosures from AI developers are critical to ensuring AI tools work as advertised and enhance practice operations.
- **Policies should be aligned across agencies to avoid establishing competing and confusing standards.** Federal regulation of AI should adequately balance the promise of AI technological capabilities along with the potential risks.
 - **The deployment of AI should avoid the unintentional exacerbation of current administrative hurdles.** Federal and private payers should not use AI to amplify burdens associated with prior authorization and intensify denials of critical patient care.
 - **Payers must be transparent and provide ample disclosures** about their use of AI for utilization management, claims processing, and coverage limitations. AI automated systems utilized by payers must be ethically designed and evidence-based, and should not interfere with physician clinical decision making.
 - **Patient privacy should remain a priority first.** Healthcare data used to develop and implement AI technology should be subject to sensible and robust security and privacy protections.
 - **All attempts should be made to mitigate discrimination and bias** in the development and utilization of AI to ensure these systems do not perpetuate harmful healthcare inequities.
 - **Medical Groups, physicians, and other providers should be appropriately protected** from liability associated with AI as it pertains to the conditions of the technology developed outside of the practice.

NATIONAL HEALTH COUNCIL

The National Health Council (NHC) appreciates the Senate Finance Committee holding this hearing today on **“Artificial Intelligence and Health Care: Promise and Pitfalls.”** This is an issue of great importance to patients and the patient community. The NHC recognizes that this hearing focuses on both the potential benefits and risks of artificial intelligence (AI) for patients, as the patient community is also acutely focused on this duality. We advance this statement to assure that the patient’s perspective is front and center in the hearing discussion and record, since none of the scheduled witnesses directly represent that viewpoint. Moving forward, the patient community expects to be included and engaged with policymakers to craft appropriate policies that assure that patients—the end users of health care—benefit from, and are protected from, the use of AI in health care. The NHC is a unique and ready resource on this critical topic.

Created by and for patient organizations over 100 years ago, the NHC brings diverse organizations together to forge consensus and drive patient-centered health policy. We promote increased access to affordable, high-value, equitable, and sustainable health care. Made up of 170 national health-related organizations and businesses, the NHC’s core membership includes the nation’s leading patient organizations. Other members include health-related associations and nonprofit organizations including the provider, research, and family caregiver communities; and businesses and organizations representing biopharmaceuticals, devices, diagnostics, generics, and payers.

Promises

Advances AI are increasingly being used to transform every facet of health care—such as improving accuracy of medical imaging and diagnoses, managing provider workflow, and speeding research and development pathways—and are having a substantial direct and indirect impact on patients. These advances hold tremendous promise to help increase the quality, timeliness, and equity of care. However, as the Committee recognizes in the title of this hearing, there is also tremendous potential for pitfalls that can harm patients, such as automating coverage denials. **To fully realize AI’s promise and minimize its pitfalls means that policy and regulatory initiatives must elevate and reflect the interests, concerns, and perspectives of patients as part of a collaborative approach.** One significant concern is the ongoing issues of developing AI that could amplify existing biases in the health care system. The data and technology used to develop and operationalize AI needs to be as free of bias as possible, otherwise existing health inequities will be further embedded in care. The NHC calls for developers, manufacturers, practitioners, patients, policymakers, regulators, and other stakeholders to engage and collaborate to continuously improve the safety and quality of AI technologies as con-

ditions evolve. **All stakeholders must be a part of the future of AI in health care, but at the very forefront must be the patient community.**

Pitfalls

AI has the potential to dramatically improve health care research, delivery, and access for patients, but only if its applications are implemented in a careful and responsible manner that accounts for and minimizes its risks. As the use of AI and other emerging technologies evolve and expand, there is a growing need to minimize potential risks which includes unintended consequences of use, adverse events, overriding patient and provider expertise, inadvertent reinforcement of implicit and explicit biases and inequities, inaccuracies in training data that lead to hard-to-detect and misleading results, and the weakening of patient privacy protections.

Sample Key Components of Responsible AI Use

AI's integration into health care delivery must be grounded in a commitment to enhancing patient access to care, advancing the quality of care, and improving operational efficiency. This must be achieved through thoughtful and effective implementation and careful and continuous oversight. The use of AI in health care decision-making must also support and supplement, not supplant, human decision-making, patient preferences, and clinician knowledge. In addition, the individuality of each patient must be recognized and supported. To achieve this, the NHC urges patient engagement in the development and operationalizing of health care tools that rely on AI to assure they reflect their needs and preferences.

Everyone, including the patient community, is continuing to learn about opportunities and challenges in leveraging AI, as its use continuously advances. This means that policies must be flexible enough to encompass new and emerging use cases while not undermining the existing policies and protections governing the health care industry. Consistent and ongoing engagement with patients will be paramount. While our perspective will evolve with these technologies, the below list demonstrates our current thinking about some of the key components and characteristics of the responsible use of AI-enabled technologies in health care from a patient perspective:

- AI applications in health care must be trustworthy, unbiased, ethical, fair, appropriate, valid, effective, safe, grounded in evidence, subject to governance controls and meaningful oversight, and safeguarded by robust privacy and security.
- AI must be used to advance health equity and not further drive health disparities.
- AI tools that will be used in health care, particularly those that are used by patients and/or directly affect patient care or coverage, must be developed with patient input into the effect of algorithms, devices, and other aspects of AI creation, use, and analysis.
- Expert human oversight of many AI uses is critical to maintaining safety and accuracy and ensuring continuous improvements to retrain as conditions change.
- Pre-deployment testing should be conducted in a diverse range of real-world clinical settings.
- Information derived from AI-enabled outputs to inform health care decision-making should:
 - Be accessible, explainable, reproducible, and understandable to the intended audience;
 - Detail the benefits and limitations of a given AI-enabled technology;
 - Have privacy and security standards for safeguarding patient information in place; and
 - Mitigate potential biases that could exacerbate health disparities and promote health equity.
- Robust and continuous feedback loops should be created, leveraged, and optimized to identify and mitigate the risk of harms.
- Users should be properly trained on intended applications, system capabilities and limitations, real-world use cases, and the probabilistic nature of AI.

Conclusion

The NHC values this opportunity to engage in this critical dialogue on AI in health care. Please do not hesitate to contact Eric Gascho, Senior Vice President of Policy and Government Affairs, if you or your staff would like to discuss these comments in greater detail. He is reachable via e-mail at egascho@nhcouncil.org.

NATIONAL HEALTH LAW PROGRAM

1444 I Street, NW, Suite 1105
 Washington, DC 20005
 (202) 289-7661
<https://healthlaw.org/>

The National Health Law Program (“NHeLP”) submits this testimony to the Senate Finance Committee regarding the use of algorithms and artificial intelligence (AI) systems in health care. NHeLP is a public interest law firm that fights for equitable access to quality health care for people with low incomes and underserved populations, and for health equity for all. For over 50 years, we have litigated to enforce health care and civil rights laws, advocated for better federal and state health laws and policies, and trained, supported, and partnered with health and civil rights advocates across the country. NHeLP’s testimony is based on our long history of advocacy to protect Medicaid beneficiaries against harmful automated decision-making systems (ADS), such as algorithms and AI.

For decades NHeLP has identified errors, discrimination, and due process violations in ADS and fought against them. We have real-world experience fighting the harm caused by technology in public benefits, knowledge about the how and why such harms occur, and practical ideas about policies necessary to protect against such harms.¹ This experience gives us a different and needed perspective on policy efforts to protect against harmful AI. We understand what the systems look like on the ground and how they impact people’s rights. As part of our work to ensure technology helps rather than harms Medicaid enrollees, NHeLP has partnered with other advocates, including tech justice advocates, to advance protections in public benefits programs so that people are not wrongfully denied benefits they need. For example, we have partnered with Upturn and Legal Aid of Arkansas to form the Benefits Tech Advocacy Hub to give advocates tools to fight harmful benefits technology and force greater transparency so that harm to individuals can be identified, prevented, or reduced earlier in the technology’s lifecycle. In addition to ongoing advocacy regarding individual ADS, NHeLP has also released our Principles for Fairer, More Responsive Automated Decision-Making Systems (“ADS Principles”), which reflect our years of work regarding ADS, including AI, and what features and protections are needed in responsible ADS.

In the past several years as interest in algorithmic accountability has grown, we consistently see proposals to mitigate the harms of ADS that fail to recognize the impact on public benefits, for which it is well-recognized that people have a “brutal need.”² Protections of notice, transparency, and explainability already exist, are constitutionally required, and must be fully recognized and enforced in any AI policies that impact public benefits.³ We welcome the opportunity to offer testimony to this Committee. NHeLP asks that this Committee:

- Use a broad definition of AI to include all of the types of AI currently causing harm to people’s access to care.
- Embrace NHeLP’s ADS Principles regarding transparency, protection of civil rights, user-focus, validity, mitigation of bias, and humility and redundancy as well as the work of the Benefits Tech Advocacy Hub in the Committee’s work on algorithms and AI systems in health care.
- Recognize that individuals receiving health care provided through a public benefit program such as Medicaid have specific rights and protections that demand greater transparency, nondiscrimination, and accountability than many AI fairness proposals include; these rights cannot be ignored in legislative approaches. We also ask that business interests such as trade secret protections not be allowed to stand in the way of transparency about the source, testing, and decision-making of AI.

AI Protections Must Include a Broad Definition of AI to Address Current ADS Harms

Automation can facilitate access to benefits and increase efficiency, but ADS that is poorly designed, based on biased research or data sets, not implemented with appropriate testing, and not adequately monitored creates significant harm. Particu-

¹ See Nat’l Health Law Program, *Fairness in Automated Decision Making Systems*, <https://healthlaw.org/algorithms/>.

² *Goldberg v. Kelly*, 397 U.S. 254 (1970).

³ *Id.*; see also Jane Perkins, Nat’l Health Law Program, *Demanding Ascertainable Standards*, Nat’l Health Law Program (June 11, 2021), <https://healthlaw.org/resource/demanding-ascertainable-standards-medicaid-as-a-case-study/>.

larly in Medicaid, this harm affects people who have very few resources to absorb it.⁴ When ADS harms Medicaid recipients, they have life altering losses of benefits; this loss of Medicaid coverage is a well-recognized harm.⁵ However, not all ADS creating harms in Medicaid meet all definitions of AI. For example, when asked to inventory AI use cases in response to Executive Order 13960 regarding the use of trustworthy AI in the federal government, the Department of Health and Human Services created a three-page list.⁶ This list does not include the Federal Marketplace for health care coverage that processed applications and plan selection for over 20 million people, with more being transferred to Medicaid, through an automated system that implements a complex system of eligibility rules, including state specific rules, to determine whether a person is eligible for health coverage or not.⁷ While the Federal Marketplace, like many ADS in Medicaid that have caused harm, is not a complex machine learning version of AI, it is the type of AI that can and has caused immense harm to people who are wrongly determined ineligible for coverage or assistance paying their premiums. Nor did HHS's list include the millions upon millions of federal dollars that have been spent building state automated eligibility systems, many of which have or have had significant issues, causing improper terminations of Medicaid coverage and harm millions of individuals.⁸ These Medicaid eligibility systems annually process the nearly 90 million people enrolled in Medicaid and CHIP throughout the country.⁹

Both the NHeLP and Benefits Tech Advocacy Hub websites include examples of harm from various types of AI in Medicaid. Wrongful decisions by AI in Medicaid have caused people to lose health care coverage for which they were eligible and not be able to fix the problem without advocacy intervention, lose eligibility for and need hours of critical home and community-based services (HCBS) that keeps people safe and healthy in their homes and out of institutions. They have also denied needed care through harmful prior authorization tools based in fiscal decisions rather than generally accepted standards of care.¹⁰ Regardless of the level of sophistication or complexity of the AI, protections must be in place. The harm from machine learning is no greater than that generated by an algorithm written by a State Medicaid agency or its contractor to determine eligibility for HCBS—both deny critical care and cause life-long harm. We ask that this Committee not be distracted by the complexity of AI such as machine learning, but recognize that any legislative action regarding AI should be broadly inclusive in order to protect against harm.

⁴ See, e.g., Sarah Grusin et al., Nat'l Health Law Program, FTC Complaint: Request for Investigation into Deloitte's Texas Medicaid Eligibility System 30–37 (Jan. 31, 2024), <https://healthlaw.org/resource/ftc-complaint-request-for-investigation-into-deloittes-texas-medicaid-eligibility-system/>.

⁵ See, e.g., *Smith v. Benson*, 703 F. Supp. 2d 1262 (S.D. Fla. 2010); Benjamin D. Sommers et al., *Health Insurance Coverage and Health—What the Recent Evidence Tells Us*, 377 NEW ENGLAND J. MED. 586 (2017), <http://www.nejm.org/doi/full/10.1056/NEJMsb1706645>; Benjamin D. Sommers, *State Medicaid Expansions and Mortality, Revisited: A Cost-Benefit Analysis*, 3 AM. J. OF HEALTH ECON. 392 (2017), <https://dash.harvard.edu/bitstream/handle/1/27305958/Mcaid%20Mortality%20Revisited%20DASH%20Version.pdf?sequence=1&isAllowed=y>; Allyson G. Hall et al., *Lapses in Medicaid Coverage: Impact on Cost and Utilization Among Individuals with Diabetes Enrolled in Medicaid*, 48 MEDIC. CARE 1219 (2008); Andrew Bindman et al., *Interruptions in Medicaid Coverage and Risk for Hospitalization for Ambulatory Care-Sensitive Conditions*, 149 ANNALS INTERNAL MED. (2008), <https://www.commonwealthfund.org/publications/journal-article/2008/dec/interruptions-medicaid-coverage-and-risk-hospitalization>; Steffie Woolhandler & David U. Himmelstein, *The Relationship of Health Insurance and Mortality: Is Lack of Insurance Deadly?*, 167 ANN. INTERN. MED. 424 (2017), <http://annals.org/aim/fullarticle/2635326/relationship-health-insurance-mortality-lack-insurance-deadly>; Aviva Aron-Dine, *Ctr. on Budget and Policy Priorities, Eligibility Restrictions in Recent Medicaid Waivers Would Cause Many Thousands of People to Become Uninsured* (Aug. 9 2018), <https://www.cbpp.org/sites/default/files/atoms/files/8-9-18health.pdf>.

⁶ U.S. Dep't of Health & Human Servs., Department of Health and Human Services: Artificial Intelligence Use Cases Inventory, <https://www.hhs.gov/about/agencies/asa/ocio/ai/use-cases/index.html>.

⁷ Ctrs. for Medicare & Medicaid Servs., *Under the Biden-Harris Administration, Over 20 Million Selected Affordable Health Coverage in ACA Marketplace Since Start of Open Enrollment Period, a Record High* (Jan. 10, 2024), <https://www.cms.gov/newsroom/press-releases/under-biden-harris-administration-over-20-million-selected-affordable-health-coverage-aca>.

⁸ See, e.g., TX FTC Complaint *supra* note 4, at 12–21.

⁹ Medicaid.gov, October 2023 Medicaid & CHIP Enrollment Data Highlights, <https://www.medicare.gov/medicaid/program-information/medicaid-and-chip-enrollment-data/report-highlights/index.html>.

¹⁰ See generally Benefits Tech Advocacy Hub, Case Studies Library, <https://www.btah.org/case-studies.html>.

Embrace NHeLP's Expertise and that of its Partners

NHeLP's long history of advocacy on ADS issues has led us to think about preventive advocacy rather than only addressing ADS after they have begun to harm individuals. Our ADS Principles and work with the Benefits Tech Advocacy Hub recognize that there are benefits to automation, but those must be realized while minimizing the drawbacks. There must be thoughtful policy interventions that address each step of the ADS lifecycle so that harm can be recognized, evaluated, and remediated. Our recent experiences, including those related to the unwinding of the Medicaid continuous coverage provisions during the public health emergency, reiterate to us that ADS are often generating harmful, yet preventable, errors.¹¹ Importantly, our work understands that ADS, even if carefully created and monitored, is likely going to have errors either because of the system or because of user interaction with the system. Therefore, we have thought through both the protections needed for the system itself and the processes around the system that should function as a safety net to prevent harm.

A key part of NHeLP's work is our relationships with advocates across the country.¹² These relationships are critical to our ADS work because they help identify systems that are proposed or are actively generating harm. Our work and that of our partners, including our tech justice partners, means we have real-world examples of problems, their impact on individuals, and solutions for preventing those harms. Our community and partners have the right mix of knowledge, including legal and technical, to identify problems and solutions that will actually work.

Preventing Harm from AI in Health Care Must Incorporate Existing Legal Rights

New AI fairness and accountability principles have been emerging over the past several years, but not all of them recognize existing legal rights in recommendations of transparency and protections against bias. And few tackle the significant barrier to transparency of trade secret and similar protections. While we recognize the business interests in technology, key legal rights of those impacted by the technology must be acknowledged in AI policy efforts. Importantly, some level of transparency is required when ADS is making decisions about public benefits due to constitutional due process requirements.¹³ This is not "optional" or a "best practice." In addition, public benefits ADS transparency may also be required by other laws, including public records.¹⁴

As described in NHeLP's ADS Principles regarding transparency, without transparency throughout the lifecycle of an ADS, problems and the harms they cause will likely come to light only when sufficient numbers of people are harmed to identify there is a problem. But even then, if transparency is not required, that the ADS is at fault and what needs to be addressed cannot occur and harm is likely to continue. For many impacted by ADS in Medicaid, once the harm has occurred, it is not easily ameliorated either because they do not readily return to the program, they have difficulty re-enrolling, or the service denial causes a domino effect of harms.¹⁵

¹¹See, e.g., Nat'l Health Law Program, Fairness in Automated Decision Making, <https://healthlaw.org/algorithms/>; TX FTC complaint, *supra* note 4.

¹²See, e.g., Nat'l Health Law Program, Health Law Partnerships, <https://healthlaw.org/health-law-partnerships/>.

¹³Perkins, *supra* note 3.

¹⁴See, e.g., *K.W. by D.W. v. Armstrong*, No. 1:12-CV-00022-BLW-CWD, 2023 WL 5431801 (D. Idaho Aug. 23, 2023) (ordering disclosure of manual information related to algorithm for services was necessary to comply with due process and did not infringe upon copyright restrictions); *Salazar v. D.C.*, 750 F. Supp. 2d 65 (D.D.C. 2010) (providing limited protected order to disclosed standards that were asserted to be protected by business interests); *Arkansas Dep't of Com., Div. of Workforce Servs. v. Legal Aid of Arkansas*, 2022 Ar. 130, 546 S.W.3d 9 (2022) (finding trade secret protections in public records did not limit beneficiary access to algorithm used in unemployment algorithm).

¹⁵Sarah Sugar et al., ASPE Office of Health Policy, *Medicaid Churning and Continuity of Care: Evidence and Policy Considerations Before and After the COVID-19 Pandemic* (Apr. 12, 2021), https://aspe.hhs.gov/sites/default/files/migrated_legacy_files/199881/medicaid-churning-ib.pdf (finding that Medicaid churn leads to periods of uninsurance, delayed care, and less preventative care for beneficiaries and higher administrative costs, less predictable state expenditures, and higher monthly care costs); Sophie Novack, *As Texas Throws 1.8 Million Off Medicaid, Children Pay the Price*, TEXAS MONTHLY (Jan. 25, 2024), <https://www.texasmonthly.com/news-politics/medicaid-disenrollment-texas-children/> (describing the impact of eligibility system requesting documentation it should already have access to and a child losing services and enrollment in treatment program critical to her walking and balance); Sarah Grusin & Eliz-

Conclusion

We ask that this committee recognize that efforts to address harm from ADS in health care must include approaches to address those harms in Medicaid and other government-funded health care programs. And that those efforts recognize not only the unique rights of enrollees, but the extent of harm as well. Our ADS Principles set forth our asks regarding algorithmic fairness and we welcome questions and conversations to further provide information and guidance on policy solutions that will provide meaningful protections to current harms.

For further information or questions about this testimony, please contact Elizabeth Edwards at the National Health Law Program by email at edwards@healthlaw.org.

NATIONAL NURSES UNITED

February 6, 2024

Senate Committee on Finance
219 Dirksen Senate Office Building
Washington, DC 20510

Dear Chairman Wyden and Ranking Member Crapo:

In light of the full committee hearing today titled “Artificial Intelligence and Health Care: Promise and Pitfalls,” I write to you on behalf of National Nurses United, the nation’s largest union and professional association of registered nurses (RNs) to discuss the ways that our nearly 225,000 members are already experiencing the impacts of artificial intelligence (AI) and data-driven technologies at the hospital bedside.

The decisions to implement AI technologies are often made without the knowledge of either nurses or patients, and are putting patients and the nurses who care for them at risk. AI technology is being used to replace educated registered nurses exercising independent judgment with lower-cost staff following algorithmic instructions. However, patients are unique and health care is made up of non-routine situations that require human touch, care, and input. **AI poses significant risks to patient care and to nursing practice, and all legislative and regulatory steps taken must utilize the precautionary principle—an idea at the center of public health analysis—in order to protect patients from harm.**

NNU urges the Federal Government to pursue a regulatory framework that safeguards the clinical judgment of nurses and other health care workers from being undermined by AI and other data-driven technologies. NNU recommends that Congress take the following actions:

All statutes and regulations must be grounded in the precautionary principle. NNU urges Congress to develop regulations that require technology developers and health care providers to prove that AI and other data-driven digital technologies are safe, effective, and therapeutic for both a specific patient population and the health care workforce engaging with these technologies before they are deployed in real-world care settings. This goes beyond racial, gender, and age-based bias. As each patient has unique traits, needs, and values, no AI can be sufficiently fine-tuned to predict the appropriate diagnostic, treatment, and prognostic for an individual patient. Liability for any patient harm associated with failures or inaccuracies of automated systems must be placed on both AI developers and health care employers and other end users. Patients must provide informed consent for the use of AI in their treatment, including notification of any clinical decision support software being used.

Privacy is paramount in health care—Congress must prohibit the collection and use of patient data without informed consent, even in so-called de-identified form. There are often sufficient data points to reidentify so-called de-identified patient information. Currently, health care AI corporations institute gag clauses on users’ public discussions of any issues or problems with their products

abeth Edwards, Nat’l Health Law Program, *Recent Filing in Lawsuit Describes Medicaid Unwinding Harms in Tennessee* (Aug. 2, 2023) (describing issues with TennCare’s eligibility system, including having to repeatedly submit the same information, not properly being found eligible, and requiring advocacy intervention); see generally Nat’l Health Law Program, *A.M.C. v. Smith*, Middle District of Tennessee, <https://healthlaw.org/resource/a-m-c-v-smith-middle-district-of-tennessee/> (case involving issues with notices, the eligibility system, disability discrimination, and access to fair hearings to address errors).

or cloak the workings of their products in claims of proprietary information. Such gag clauses must be prohibited by law. Additionally, health care AI corporations and the health care employers that use their products regularly claim that clinicians' right to override software recommendations makes them liable for any patient harm while limiting their ability to fully understand and determine how they are used. Thus, clinicians must have the legal right to override AI. For nurses, this means the right to determine nurse staffing and patient care based on our professional judgment.

Patients' informed consent and the right to clinician override are not sufficient protections, however. **Nurses must have the legal right to bargain over the employer's decision to implement AI and over the deployment and effects of implementation of AI in our workplace.** In addition to statutes and regulations codifying nurses' and patients' rights directly, Congress needs to strengthen workers' rights to organize, collectively bargain, and engage in collective action overall. Health care workers should not be displaced or deskilled as this will inevitably come at the expense of both patients and workers. At the regulatory level, the Centers for Medicare and Medicaid Services must require health care employers to bargain over any implementation of AI with labor unions representing workers as a condition of participation.

Congress must protect workers from AI surveillance and data mining. Congress must prohibit monitoring or data mining of worker-owned devices. Constant surveillance can violate an employee's personal privacy and personal time. It can also allow management to monitor union activity, such as conversations with union representatives or organizing discussions, which chills union activity and the ability of workers to push back against dangerous management practices. The federal government must require that employers make clear the capabilities of this technology and provide an explanation of how it can be used to track and monitor nurses. Additionally, Congress must prohibit the monitoring of worker location, data, or activities during off time in devices used or provided by the employer. Employers should be restricted from collecting biometric data or data related to workers' mental or emotional states. Finally, employers should be prohibited from disciplining an employee based on data gathered through AI surveillance or data mining, and AI developers and employers should also be prohibited from selling worker data to third parties.

National Nurses United looks forward to future conversations on this topic, and to working with this committee to ensure that the federal government develops effective regulations that will protect nurses and patients from the harm that can be caused by artificial intelligence and data-driven technologies in health care.

Sincerely,

Amirah Sequeira
National Government Relations Director

