

**GAME CHANGERS: ARTIFICIAL INTELLIGENCE
PART III, ARTIFICIAL INTELLIGENCE AND
PUBLIC POLICY**

HEARING
BEFORE THE
SUBCOMMITTEE ON
INFORMATION TECHNOLOGY
OF THE
COMMITTEE ON OVERSIGHT
AND GOVERNMENT REFORM
HOUSE OF REPRESENTATIVES
ONE HUNDRED FIFTEENTH CONGRESS
SECOND SESSION

APRIL 18, 2018

Serial No. 115-79

Printed for the use of the Committee on Oversight and Government Reform



Available via the World Wide Web: <http://www.fdsys.gov>
<http://oversight.house.gov>

U.S. GOVERNMENT PUBLISHING OFFICE

31-118 PDF

WASHINGTON : 2018

COMMITTEE ON OVERSIGHT AND GOVERNMENT REFORM

Trey Gowdy, South Carolina, *Chairman*

John J. Duncan, Jr., Tennessee	Elijah E. Cummings, Maryland, <i>Ranking</i>
Darrell E. Issa, California	<i>Minority Member</i>
Jim Jordan, Ohio	Carolyn B. Maloney, New York
Mark Sanford, South Carolina	Eleanor Holmes Norton, District of Columbia
Justin Amash, Michigan	Wm. Lacy Clay, Missouri
Paul A. Gosar, Arizona	Stephen F. Lynch, Massachusetts
Scott DesJarlais, Tennessee	Jim Cooper, Tennessee
Virginia Foxx, North Carolina	Gerald E. Connolly, Virginia
Thomas Massie, Kentucky	Robin L. Kelly, Illinois
Mark Meadows, North Carolina	Brenda L. Lawrence, Michigan
Ron DeSantis, Florida	Bonnie Watson Coleman, New Jersey
Dennis A. Ross, Florida	Raja Krishnamoorthi, Illinois
Mark Walker, North Carolina	Jamie Raskin, Maryland
Rod Blum, Iowa	Jimmy Gomez, Maryland
Jody B. Hice, Georgia	Peter Welch, Vermont
Steve Russell, Oklahoma	Matt Cartwright, Pennsylvania
Glenn Grothman, Wisconsin	Mark DeSaulnier, California
Will Hurd, Texas	Stacey E. Plaskett, Virgin Islands
Gary J. Palmer, Alabama	John P. Sarbanes, Maryland
James Comer, Kentucky	
Paul Mitchell, Michigan	
Greg Gianforte, Montana	

SHERIA CLARKE, *Staff Director*

WILLIAM MCKENNA, *General Counsel*

TROY STOCK, *Information Technology Subcommittee Staff Director*

SARAH MOXLEY, *Senior Professional Staff Member*

SHARON CASEY, *Deputy Chief Clerk*

DAVID RAPALLO, *Minority Staff Director*

SUBCOMMITTEE ON INFORMATION TECHNOLOGY

Will Hurd, Texas, *Chairman*

Paul Mitchell, Michigan, <i>Vice Chair</i>	Robin L. Kelly, Illinois, <i>Ranking Minority</i>
Darrell E. Issa, California	<i>Member</i>
Justin Amash, Michigan	Jamie Raskin, Maryland
Steve Russell, Oklahoma	Stephen F. Lynch, Massachusetts
Greg Gianforte, Montana	Gerald E. Connolly, Virginia
	Raja Krishnamoorthi, Illinois

CONTENTS

Hearing held on April 18, 2018	Page 1
WITNESSES	
Mr. Gary Shapiro, President, Consumer Technology Association	
Oral Statement	4
Written Statement	6
Mr. Jack Clark, Director, OpenAI	
Oral Statement	12
Written Statement	14
Ms. Terah Lyons, Executive Director, Partnership on AI	
Oral Statement	23
Written Statement	25
Dr. Ben Buchanan, Postdoctoral Fellow, Cyber Security Project, Science, Technology, and Public Policy Program, Belfer Center for Science and Inter- national Affairs, Harvard Kennedy School	
Oral Statement	37
Written Statement	39

GAME CHANGERS: ARTIFICIAL INTELLIGENCE PART III, ARTIFICIAL INTELLIGENCE AND PUBLIC POLICY

Wednesday, April 18, 2018

HOUSE OF REPRESENTATIVES,
SUBCOMMITTEE ON INFORMATION TECHNOLOGY,
COMMITTEE ON OVERSIGHT AND GOVERNMENT REFORM,
Washington, D.C.

The subcommittee met, pursuant to call, at 2:02 p.m., in Room 2154, Rayburn House Office Building, Hon. Will Hurd [chairman of the subcommittee] presiding.

Present: Representatives Hurd, Issa, Amash, Kelly, Connolly, and Krishnamoorthi.

Mr. HURD. The Subcommittee on Information Technology will come to order. And without objection, the chair is authorized to declare a recess at any time.

Good afternoon. I welcome y'all to our final hearing in our series on artificial intelligence. I've learned quite a bit from our previous two hearings, and I expect today's hearing is going to be equally informative.

This afternoon, we are going to discuss the appropriate roles for the public and private sectors as AI, artificial intelligence, matures.

AI presents a wealth of opportunities to impact our world in a positive way. For those who are vision impaired, there is AI that describes the physical world around them to help them navigate, making them more independent. AI helps oncologists target cancer treatment more quickly. AI has the potential to improve government systems so that people spend less time trying to fix problems, like Social Security cards or in line at Customs.

As with anything that brings tremendous potential for rewards, there are great challenges ahead as well. AI can create video clips of people saying things they did not say and would never support. AI tools and cyber attacks can increase the magnitude and reach of those—of these attacks to disastrous levels.

In addition, both our allies and potential adversaries are pursuing AI dominance. It is not a foregone conclusion that the U.S. will lead in this technology. We need to take active steps to ensure America continues to be the world leader in AI.

On the home front, bias, privacy, ethics, and the future of work are all challenges that are a part of AI. So given the great possibilities and equal great potential hardships, what do we do? What is the role of government in stewarding this great challenge to benefit

all? What should the private sector be doing to enhance the opportunity to minimize the risk?

While I do not expect anyone to have all these answers today, I think our panel of witnesses will have suggestions for the way forward when it comes to AI.

While this is the final hearing in our AI series, this work does not end today. And our subcommittee will be releasing a summary of what we have learned from the series in the coming weeks outlining steps we believe should be taken in order to help drive AI forward in a way that benefits consumers, the government, industry, and most importantly, our citizens.

I thank the witnesses for being here today and look forward to learning from y'all, and we can all benefit from the revolutionary opportunities AI offers. And as always, I'm honored to be exploring these issues in a bipartisan fashion with my friend, the ranking member, the woman, the myth, the legend, Robin Kelly from the great State of Illinois.

Ms. KELLY. Thank you so much.

Thank you, Chairman Hurd, and welcome to all of our witnesses here today. This is the third hearing, as you've heard, that our subcommittee has held on the important topic of artificial intelligence, or AI. Our two prior hearings have shown how critical the collection of data is to the development and expansion of AI. However, AI's reliance on the use of personal information raises legitimate concerns about personal privacy.

Smart devices of all kinds are collecting your data. Many of us have to look no further than the smart watch on our wrists to see this evidence in motion. The arms race to produce individual predictive results is only increasing with smart assistants like Alexa and Siri in your pocket and listening at home for your next command. Sophisticated algorithms help these machines refine their suggestions and place the most relevant information in front of our customers.

These systems, however, rely upon vast amounts of data to produce precise results. Privacy concerns for tens of millions of Facebook users were triggered when the public learned that Cambridge Analytica improperly obtained to potentially use their personal data to promote the candidacy of Donald Trump.

Whether Congress passes new laws or industry adopts new practices, clearly, consumers need and deserve new protections. To help us understand what some of these protections may look like, Dr. Ben Buchanan from Harvard University's Belfer Center for Science and International Affairs, is here with us today. Dr. Buchanan has written extensively on the different types of safeguards that may be deployed on AI systems to protect the personal data of consumers.

Advancement in AI also pose new challenges to cybersecurity due to increased risk of data breaches by sophisticated hackers. Since 2013, we have witnessed a steady increase in the number of devastating cyber attacks against both the private and the public sectors. This past September, Equifax announced that hackers were able to exploit a vulnerability on their systems, and as a result, gained access to the personal data of over 140 million Americans.

A recent report coauthored by OpenAI, represented by Mr. Clark today, expressly warns about the increased cyber risks the country faces due to AI's advancements. According to the report, continuing AI advancements are likely to result in cyber attacks that are, quote, "more effective, more finely targeted, more difficult to attribute, and more likely to exploit vulnerabilities in AI systems."

As AI advances, another critical concern is its potential impact on employment. Last year, the McKinsey Global Institute released the findings from a study on the potential impact of AI-driven automation on jobs. According to the report, and I quote, "Up to one-third of the workforce in the United States and Germany may need to find work in new occupations."

Other studies indicate that the impact on U.S. workers may even be higher. In 2013, Oxford University reported on a study that found that due to AI automation, I quote, "about 47 percent of total U.S. employment is at risk."

To ensure that AI's economic benefits are more broadly shared by U.S. workers, Congress should begin to examine and develop policies and legislation that would assist workers whose jobs may be adversely affected by AI-driven automation.

As AI advances continue to develop, I'll be focused on how the private sector, Congress, and regulators can work to ensure that consumers' personal privacy is adequately protected and that more is being done to account for the technology's impact on cybersecurity and our economy.

I want to thank our witnesses again for testifying today, and I look forward to hearing your thoughts on how we can achieve this goal.

And again, thank you, Mr. Chairman.

Mr. HURD. I appreciate the ranking member.

And now, it's a pleasure to introduce our witnesses. Our first guest is known to everyone who knows anything about technology, Mr. Gary Shapiro, president of the Consumer Technology Association. Thanks for being here.

Mr. Jack Clark is here as well, director at OpenAI.

We have Ms. Terah Lyons, the executive director at Partnership on AI.

And last but not least, Dr. Ben Buchanan, postdoctoral fellow at Harvard Kennedy School's Belfer Center for Science and International Affairs. Say that three times fast.

I appreciate all y'all's written statements. It really was helpful in understanding this issue.

And pursuant to committee rules, all witnesses will be sworn in before you testify, so please stand and raise your right hand.

Do you solemnly swear or affirm that you're about to tell the truth, the whole truth, and nothing but the truth so help you God?

Thank you. Please be seated.

Please let the record reflect that all witnesses answered in the affirmative.

And now, in order to allow for time for discussion, please limit your testimony to 5 minutes. Your entire written statement will be made part of the record.

As a reminder, the clock in front of you shows your remaining time; the light turns yellow when you have 30 seconds left; and

when it's flashing red, that means your time is up. Also, please remember to push the talk button to turn your microphone on and off.

And now, it's a pleasure to recognize Mr. Shapiro for your opening remarks.

WITNESS STATEMENTS

STATEMENT OF GARY SHAPIRO

Mr. SHAPIRO. I'm Gary Shapiro, president and CEO of the Consumer Technology Association, and I want to thank you, Chairman Hurd and Ranking Member Kelly, for inviting me to testify on this very important issue, artificial intelligence.

Our association represents 2,200 American companies in the consumer technology industry. We also own and produce the coolest, greatest, funnest, most important, and largest business and innovation event in the world, the CES, held each January in Las Vegas.

Our members develop products and services that create jobs. They grow the economy and they improve lives. And many of the most exciting products coming to market today are AI products.

CTA and our member companies want to work with you to figure out how we can ensure that the U.S. retains its position as the global leader in AI, while also proactively addressing the pressing challenges that you've already raised today.

Last month, we released a report on the current and future prospects of AI, and we found that AI will change the future of everything, from healthcare and transportation to entertainment security. But it will also raise questions about jobs, bias, and cybersecurity. We hope our research, along with the efforts of our member-driven artificial intelligence working group, will lay the groundwork for policies that will foster AI development and address the challenges AI may create.

First, consider how AI is creating efficiency and improving lives. The U.S. will spend \$3.5 trillion on healthcare this year. The Federal Government shoulders over 28 percent of that cost. By 2047, the CBO estimates Federal spending for people age 65 and older who receive Social Security, Medicare, and Medicaid benefits could account for almost half of all Federal spending.

AI can be part of the solution. Each patient generates millions of data points every day, but most doctors' offices and hospitals are not now maximizing the value of that data. AI can quickly sift through and identify aspects of that data that can save lives. For example, Qualcomm's alert watch AI system, which provides real-time analysis of patient data during surgery, significantly lowers patients' heart attacks and kidney failures, and it reduces average hospital stays by a full day.

Cybersecurity is another area where AI can make a big impact, according to our study. AI technologies can interpret vast quantities of data to prepare better for and protect against cybersecurity threats. In fact, our report found that detecting and deterring security intrusions was a top area where companies are today using AI. AI should contribute over \$15 trillion to the global economy by 2030, according to PWC.

Both the present and prior administrations have recognized the importance of prioritizing AI. But AI is also capturing the attention of other countries. Last year, China laid out a plan to create \$150 billion world leading AI industry by 2030. Earlier this year, China announced a \$2 billion AI research park in Beijing. France just unveiled a high-profile plan to foster AI development in France and across the European Union. I was there last week, and it was the talk of France.

Today, the U.S. is the leader in AI, both in terms of research and commercialization. But as you said, Mr. Chairman, our position is not guaranteed. We need to stay several steps ahead. Leadership from the private sector, supported by a qualified talent pool and light touch regulation, is a winning formula for innovation in America. We need government to think strategically about creating a regulatory environment that encourages innovation in AI to thrive, while also addressing the disruptions we've been talking about.

Above all, as we noted in our AI report, government policies around AI need to be both flexible and adaptive. Industry and government also need to collaborate to address the impact AI is having and will have on our workforce. The truth is most jobs will be improved by AI, but many new jobs will be created and, of course, some will be lost.

We need to ensure that our workforce is prepared for these jobs of the future, and that means helping people whose jobs are displaced gain the skills that they need to succeed in new ones.

CTA's AI working group is helping to address these workforce challenges. We just hired our first vice president of U.S. jobs; and on Monday, we launched CTA's 21st Century Workforce Council to bring together leaders in our industry to address the significant skills gap in our workforce we face today.

In addition to closing the skills gap, we need to use the skills of every American to succeed. CTA is committed to strengthening the diversity of the tech workforce. Full representation of a workforce will go a long way to making sure that tech products and services consider the needs and viewpoints of diverse users.

We as an industry also need to address data security, and we also need to welcome the opportunity to continue to work with you on that in other areas. We believe that the trade agenda, the IP agenda, and immigration all tie into our success as well in AI.

There's no one policy decision or government action that will guarantee our leadership in AI, but we are confident we can work together on policies that will put us in the best possible position to lead the world in AI and deliver the innovative technologies that will change our lives for the better.

[Prepared statement of Mr. Shapiro follows:]

House Oversight Committee, Subcommittee on Information Technology hearing- “Game Changers: Artificial Intelligence Part III, Artificial Intelligence and Public Policy.”

Hearing Testimony

Testimony of Gary Shapiro, president and CEO, Consumer Technology Association (CTA)TM

Thank you Chairman Hurd, Ranking Member Kelly, and members of the subcommittee, for inviting me to testify today on the future of artificial intelligence technology. I am Gary Shapiro, president and CEO of the Consumer Technology Association.

The Consumer Technology Association is the trade association representing the \$351 billion U.S. consumer technology industry, which supports more than 15 million U.S. jobs. We own and produce CES® - the Global Stage for Innovation, held each January in Las Vegas. I am fortunate to have a front row seat each day as our members develop and introduce innovative and life-changing products and services, create jobs, and grow the economy. At CTA, we work to support public policy that fosters innovation, advances competitiveness and promotes jobs and business creation.

I'd like to thank the subcommittee for holding this three-part series of hearings on artificial intelligence (AI). AI technology is still in its early days, but it is making a significant impact on our world. Medical professionals are able to detect health problems earlier and more accurately with the help of AI. Our defense agencies are using AI to scan large amounts of data for serious threats. Millions of Americans experience AI every day in the form of voice assistants, credit card fraud detection warnings and personalized online shopping recommendations.

As AI becomes capable of doing more complex tasks, it will revolutionize even more aspects of our lives. It will also raise an increasing number of questions about jobs, bias, cybersecurity and other serious issues. CTA applauds Chairman Hurd and Ranking Member Kelly for their early recognition that AI can play a major role in improving society and making government more efficient. We also appreciate your thoughtful approach to tackling the big questions AI is already raising.

CTA and our member companies want to work with you to figure out how we can ensure the U.S. retains its position as the global leader in AI, while also proactively addressing the pressing challenges raised by this new technology. That is why we created CTA's Artificial Intelligence Working Group earlier this year. The working group includes the top companies in AI development and deployment, as well as innovative startups in the area. In the near future, our group will focus on developing AI policy principles for the industry. The group will also work on establishing uniform definitions for terms like artificial intelligence and machine learning.

Recognizing the potential of artificial intelligence, CTA worked with Nasdaq to develop the Nasdaq CTA Artificial Intelligence & Robotics (NQROBO) index. The index began on December 18, 2017. On February 21, 2018, First Trust launched the first exchange traded fund (trading under the ticker ROBT) containing the index's underlying funds.

In March of this year CTA released a report on the current and future prospects of artificial intelligence. Our report includes findings from 20 in-depth interviews with experts on AI development and applications within a broad range of industries including automotive manufacturing, health care, and retail. We explored how AI is poised to revolutionize these and other industries. We also delved in to the various challenges AI may face. I hope this report can lay the groundwork for working with members of this subcommittee on smart AI policies.

The United States has been the indisputable world leader in disruptive innovations in the internet age. Countries around the world have tried to replicate the economic success and global influence of the U.S. tech sector, but none has succeeded on a significant scale. Artificial intelligence is the next frontier in this global race. Nations recognize that whoever leads in AI will have a major advantage economically, militarily and societally. Global leadership depends on leadership in artificial intelligence.

Narrow AI, which is capable of a limited set of tasks, is already having a major impact. In CTA's recently released AI report, industry experts identified health care as one of the areas where AI could in the near future solve real world problems. Each patient generates millions of data points every day (over 21 million if they are connected to 10 devices in the operating room or ICU), and yet relatively little is typically done with that data. Less than 0.1 percent of it is recorded in a patient's electronic medical record. Even if doctors wanted to put more of that data to use, a human would not be able to analyze that volume of information in real time. AI, on the other hand, can quickly sift through and identify pertinent aspects of that data, applying predictive and prescriptive analytics to save lives.

When the Qualcomm AlertWatch AI system, which provides real-time analysis of patient data to anesthesiologists during surgery, was tested at Michigan Medicine, the University of Michigan's health system, heart attacks and kidney failure in surgery patients were significantly lowered. Average stays in the hospital were reduced by a full day.

AI's benefit to health care can be measured in the millions of lives it will potentially save and the billions of dollars of cost savings it could likely bring. Imagine the impact if doctors across the whole country were able to use AI to simultaneously improve health outcomes and reduce hospital stays by an average of one day.

The U.S. will spend \$3.5 trillion on health care this year. That number is expected to hit \$5.7 trillion by 2026 – a 63 percent increase. The federal government shoulders over 28 percent of that cost, and state and local governments are responsible for nearly 17 percent of it. Last year, the Congressional Budget Office said, "By 2047, under current law, federal spending for people age 65 and older who receive benefits from Social Security, Medicare, and Medicaid would account for about half of all federal non interest spending, compared with about two-fifths today." Getting our country on sound financial footing, particularly with regard to health care spending, is a matter of national security and stability. AI can be a key part of the solution.

Artificial intelligence also is providing solutions to many of our transportation and mobility challenges, and it is poised to power the smart cities of the future. Municipal entities, from Pima County, AZ, to the Massachusetts Department of Technology, are actively pursuing AI systems to optimize traffic flow and ease congestion. Self-driving vehicles, which rest on a foundation of AI software, may eliminate more than 90 percent of accidents caused by human error. Personal mobility solutions such as Honda's UNI-CUB, which assists people who are not able to walk long distances, use AI to detect subtle movements of their users to adjust speed and direction. Ridesharing apps including Uber and Lyft use AI algorithms to ensure drivers get to neighborhoods where and when riders need them.

Our world-leading agricultural industry is relying more and more on artificial intelligence. By analyzing soil, environmental and weather data, AI is providing farmers with powerful new tools to increase agricultural yields. On today's farms, dairy cows now can wear so-called "cow Fitbits" that monitor their movements, health and diet.

Cybersecurity is identified in CTA's AI report as another area where AI can make a big impact. AI technologies can interpret vast quantities of data to better prepare for, and protect against cybersecurity threats. Last year, when the WannaCry ransomware attacked computers across the world, AI-based solutions were able to fight back. Software that relied on Nvidia GPU-powered deep learning had extraordinary success in preventing execution of the ransomware on customers' systems.

Organizations and companies around the world are using AI to address pressing humanitarian and environmental challenges. Intel has partnered with the Center for Missing and Exploited Children to use AI to analyze the over eight million reports that the center receives every year. Human analysts could not efficiently process all this data, but with the help of AI they are able to connect relevant law enforcement agencies with the appropriate cases. The UN World Food Program is actively investigating how it can use AI to interpret data from drones and other equipment on the scene immediately after a disaster, and put that data to work in real time to inform its response. Google is using AI technology to analyze publicly available data on the location of ships to identify fisherman that may be flouting regulations.

The U.S. government is also placing a stronger emphasis on AI. As the subcommittee heard in your previous hearing on artificial intelligence, "AI part II, Artificial Intelligence and the Federal Government," numerous federal agencies are deploying AI technology to make government more effective and efficient. The Department of Homeland Security is using AI to help defend our infrastructure from the constant threat of cyberattacks. Our defense agencies are prioritizing AI development, as they recognize the significant strategic advantages it will provide. The National Science Foundation invests over \$100 million annually in foundational research related to advancements in AI. Both the Obama and Trump administrations have recognized the importance of prioritizing AI. In particular, that sentiment has been reflected strongly by the work of the Office of Science and Technology Policy under both presidents.

Artificial intelligence is expected to contribute over \$15 trillion to the global economy by 2030, according to a report by PwC. That economic figure, combined with the military and national security advantage that AI could bring, is driving significant interest in artificial intelligence

among governments around the world. Last year, China laid out its “New Generation Artificial Intelligence Plan” with a goal of creating a \$150 billion world-leading AI industry by 2030. Earlier this year, a government-funded \$2.12 billion AI research park was announced in Beijing. The Chinese plan has spawned a narrative that the U.S. is falling behind in AI.

In truth, the United States is the world leader in AI, both in terms of research and commercialization. But China’s plan to capture global leadership in the area must serve as a reminder that our position is not guaranteed. In fact, China’s plan for AI leadership in many ways mirrors the recommendations of a series of 2016 AI reports issued by the Obama administration. We need to stay several steps ahead.

China and the United States are far from the only countries that are making a big play on artificial intelligence. The French government recently unveiled a high-profile plan to foster AI development in France and the EU. In an interview with *WIRED* about the announcement, French President Emmanuel Macron highlighted innovations that he saw when he visited CES in 2016 as inspiration for France’s emphasis on developing AI in the health care space. Referring to AI, Macron was also quoted as saying, “I want to frame the discussion at a global scale. The key driver should not only be technological progress, but human progress.” Nations are not just competing to lead in the creation of cutting-edge AI technology, they are also competing to lead the conversation around how the technology will be used and how it will be regulated.

Leadership from the private sector, supported by government research and light-touch regulation, has historically served as a winning formula for innovation in America. That formula brought us the internet and America’s dominant global position in the tech industry. That same formula can be used to cement our position as the global leader in AI well into the future. We do not need a high-profile, top-down initiative from the federal government in order to lead in AI. We do, however, need the government to think strategically about creating a regulatory environment that will both encourage innovation and address the potential disruptions AI could cause. Industry experts interviewed for CTA’s AI report emphasized the fact that any government policies around AI need to be flexible and adaptive.

In addition to working together to create a pro-innovation regulatory approach, industry and government need to collaborate in tackling the serious questions raised by the increasing adoption of AI. There is understandable anxiety about the effect artificial intelligence will have on our workforce. A recent report from the Organization for Economic Cooperation and Development (OECD) found that 14 percent of jobs in OECD countries, including the U.S., are “highly automatable.” In the United States, the report estimates that 13 million jobs are at risk from automation. These numbers are lower than some other studies have suggested, but they are by no means insignificant. We need to ensure that our workforce is prepared for the jobs of the future. And we need to have plans in place to provide people whose jobs are displaced by AI and automation with the skills they need to succeed in new roles.

While some jobs will be lost due to technological advancements, many more new jobs will be created. The nature of some existing jobs will change, requiring workers to switch to new occupations, or upgrade their skills and perform new tasks. Jobs will be improved by technology – creating safer environments and requiring more social engagement and cognitive skills such as

logical reasoning and creative thinking. The recent OECD report estimated that 32 percent of jobs would be different than they are today due to technology. AI is predicted to create millions of new jobs unheard of today. People entering the workforce in nearly all sectors of our economy will need to have skill sets necessary to work alongside technology and adapt to the new job opportunities that it will bring. Many people in the workforce today will need to acquire these skills in the coming years to remain employed.

To address these workforce challenges head-on, CTA recently hired our first Vice President of US Jobs, who is staffing CTA's 21st Century Workforce Council. The council will serve as a leadership forum to address the nation's skills gap, ensure the U.S. tech sector has the high-skilled workers it needs, and devise strategies to upskill U.S. workers to succeed in the 21st century. It includes companies like Toyota, Sprint, HP, Panasonic, Sony Electronics, Bosch and others. On Monday, April 16, we briefed the White House on the council – CTA's newest member group – and on how the tech sector helps Americans prepare for the future of work.

Beyond closing the skills gap, CTA is committed to significantly improving the diversity of the tech industry workforce. This is a strategic necessity – we cannot afford to leave a significant portion of the U.S. workforce untapped. It is also a technological necessity – full representation in our workforce will go a long way toward driving innovation and ensuring tech products represent the unique needs and viewpoints of all of users. It will be much less likely for an algorithm to contain inadvertent bias if a truly diverse team is building that algorithm. As artificial intelligence emerges in high stakes environments, like criminal justice, creating diverse and inclusive work environments takes on an even greater urgency. CTA considers a diverse workforce to be a key measure of innovation, and for that reason it is one of the categories that we measure in our International Innovation Scorecard.

Another critical factor for winning public confidence in artificial intelligence is data security. AI needs a large amount of data in order to be trained for any given task. Without data, AI does not work. Because users will not embrace AI unless they know their personal data is protected, our industry needs to make sure data used in AI is secure. CTA has been a leader in creating private sector voluntary guidelines for privacy, most recently with regard to personal wellness data, and we would welcome the opportunity to continue that work in other areas. Government can help by using and promoting strong cryptography and other security protocols.

We look forward to collaborating with members of this subcommittee, and other leaders throughout Congress, to create an AI innovation agenda. This agenda should focus on AI-specific issues such as educating the future workforce for AI-related jobs, reskilling existing workers, diversity in AI development, collaboration between industry and government on critical research, and making sure AI isn't singled out for unnecessary over-regulation. It should also look at broader policy areas that could hamper or encourage American leadership in AI. Protectionist trade policies could put a damper on the global supply network for critical components and unnecessarily drive up the cost for AI-powered consumer goods. Immigration policies that allow the world's best and brightest to work and build companies in the U.S., particularly when they are educated at American universities, will help provide access to the talent we need to lead in AI. Reforming patent laws to discourage frivolous litigation would help companies funnel more money to R&D budgets and help innovative startups get off the ground.

There is no single policy decision or government action that will guarantee America's leadership in artificial intelligence. But we are confident that we can work together on a broad range of policies that will put us in the best possible position to lead well into the future and deliver the innovative technologies that will continue to change our lives for the better.

Mr. HURD. Thank you, Mr. Shapiro.
 Mr. Clark, you're now recognized for 5 minutes.

STATEMENT OF JACK CLARK

Mr. CLARK. Chairman Hurd, Ranking Member Kelly, and other members of the subcommittee, thank you for having this hearing.

I'm Jack Clark, the strategy and communications director for OpenAI. We're an organization dedicated to ensuring that powerful artificial intelligence systems benefit all of humanity. We're based in San Francisco, California, and we conduct fundamental technical research with frontiers of AI, as well as participating in the global policy discussion. I also help maintain the AI index and AI measurement and forecasting initiative which is linked to Stanford University.

I'm here to talk about how government can support AI in America, and I'll focus on some key areas. My key areas are ethics, workforce issues, and measurement. I believe these are all areas where investment and action by government will help to increase this country's chances of benefiting from this transformative technology.

First, ethics. We must develop a broad set of ethical norms governing the use of this technology, as I believe existing regulatory tools are and will be insufficient. The technology is simply developing far too quickly. As I and my colleagues recently wrote in a report *Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation*, this unprecedentedly rampant proliferation of powerful technological capabilities brings about unique threats or worsens existing ones. And because of the nature for technology, traditional arms control regimes or other policy tools are insufficient, so we need to think creatively here.

So how can we control this technology without stifling innovation? I think we need to work on norms. And what I mean by norms are developing a global sense of what is right and wrong to do with this technology. So it's not just about working here in America; it's about taking a leadership position on norm creation so that we can also influence how AI is developed worldwide. And that's something that I think the United States is almost uniquely placed to do well.

This could include new norms around publication, as well as norms around safety research, or having researchers evaluate technologies for their downsides as well as upsides and having that be a part of the public discussion.

I'm confident this will work. We've already seen similar work being undertaken by the AI community to deal with our own issues of diversity and bias. Here, norms have become a product of everyone, and by having an inclusive conversation that's involved a wide set of stakeholders, we've been able to come to solutions that don't require specific regulations but can create norms that condition the way that the innovation occurs.

So a question I have for you is, you know, what do you want to know about, what are you concerned about, and what conversations can we have to make sure that we are responsive to those concerns as a community?

Second, workforce. The U.S. is currently the leader in AI technology, but as my colleague Gary said, that's not exactly guaranteed. There's a lot of work that we need to do to ensure that that leadership remains in place, and that ranges from investment in basic research to also supporting the community of global individuals that develop AI. I mean, part of the reason we're sitting here today is because of innovations that occurred maybe 10 to 15 years ago as a consequence of people I could count on these two hands. So even losing a single individual is a deep and real problem, and we should do our best to avoid it.

Third, measurement. Now, measurement may not sound hugely flashy or exciting, but I think it actually has a lot of value and is an area where government can have an almost unique enabling role in helping innovation. You know, the reason why we're here is we want to understand AI and its impact on society, and while hearings like this are very, very useful, we need something larger. We need a kind of measurement moonshot so that we can understand where the technology is developing, you know, where it's going in the future, where new threats and opportunities are going to come from so that we can have, not only informed policymakers, but also a more informed citizenry. And I think that having citizens feel that the government knows what's going on with AI and is taking a leadership role in measuring AI's progress and articulating that back to them can make it feel like a collective across-America effort to develop this technology responsibly and benefit from it.

Some specific examples already abound for ways this works. You know, DARPA wanted to measure how good self-driving cars were and held a number of competitions, which enabled the self-driving car industry. Two years ago, it held similar competitions for cyber defense and offense, which has given us a better sense of what this technology means there. And even more recently, DIUx released their xView satellite datasets in competition, which is driving innovation in AI research in that critical area to national security and doing it in a way that's inclusive of as many smart people as possible.

So thank you very much. I look forward to your questions.

[Prepared statement of Mr. Clark follows:]

**Written Testimony of Jack Clark
Strategy and Communications Director
OpenAI**

**House of Representatives Oversight & Government Reform Committee
Subcommittee on Information Technology
Hearing on
Game Changers: Artificial Intelligence Part III – AI and Public Policy**

April 18, 2018

Chairman Hurd, Ranking Member Kelly, and Members of the Subcommittee, thank you for the opportunity to discuss the crucial subject of artificial intelligence and policy. Today's hearing comes at a critical time for the development of AI: we have a dramatic future ahead of us, filled with opportunities and challenges. I'm going to concentrate on three main areas for this hearing: ethics, workforce issues, and the importance of AI measurement and forecasting. I think these are areas where success will translate into ensuring the US remains a globally competitive place to invent and apply AI. I see this as the start of an important conversation, and one which I hope continues.

Introduction:

First, a bit about me: I work as the Strategy and Communications Director for OpenAI, a non-profit artificial intelligence research company whose goal is to ensure that powerful artificial intelligence benefits all of humanity - both through direct technical work and through analysis of its impacts.

Much of the research OpenAI does today is about pushing the frontiers of AI capabilities in specific areas and our contributions have so far included new algorithms and associated code, new tools used by other researchers, public policy work about some of the challenges and opportunities of the technology, demonstrations of the capability of powerful AI systems via our 'Dota 2' project where we have developed AI systems capable of out-competing human champions at a complex and widely played strategy game, and AI safety which is essentially the study of how to ensure that increasingly powerful systems will continue to be predictable and interpretable in what they're doing and how they are doing it while acting with greater degrees of autonomy.

I also help produce the AI Index, an AI measurement and forecasting project that is part of the Stanford One Hundred Year Study on AI. Working on AI measurement has given me some insights as to the critical importance of the measurement and forecasting of disruptive technologies and has significantly influenced my thinking about ways the government could be more involved in this critical area.

Technical Background:

it's important to clarify why we're paying attention to AI: AI's capabilities also define AI's threats and opportunities, so the rate of progress of AI conditions the environment in which we make policy decisions. As a quick refresher: in 2012 several researchers, including the co-founder of OpenAI Ilya Sutskever, showed that they could use a technology called a neural network to obtain unprecedented results on a widely-studied image recognition task. At that time, their system was able to correctly identify the objects in an image about 85 percent of the time. In the five years that have elapsed since that paper was published, accuracies have climbed to around 98 percent - surpassing the performance of humans that evaluate themselves on this task. We've seen similar trends in speech recognition as well. These innovations aren't limited to rote classification tasks - similar advances have occurred in the area of AI-aided synthesis, with new techniques invented in 2014 leading to unprecedented advances in the ability for computers to create fake images, fake audio and fake videos with levels of fidelity that approach photorealism. All of these advances have relied on the same essential machinery, much of which is available in the public domain and which runs on widely available types of computer hardware. This means progress, if anything, is set to accelerate in the coming years.

Safety and Reliability:

These new capabilities have their own unique flaws which continue to befuddle and challenge researchers, so we must remember that while this technology is capable of amazing feats of 'intelligence' it is also capable of making mistakes that seem alien to its human developers and may limit its deployment in domains that require ironclad guarantees of reliability and predictability, or may lead to the technology causing accidents when deployed. Further research in the field of AI safety (of which OpenAI has made a substantial investment) may deal with (some of) these problems.

Given that, I'll use the rest of this testimony to discuss three key areas with specific policy recommendations where success will let the US lead development of AI, which will let the US coordinate the global response to the challenges and opportunities of the technology. I'll finish my testimony by providing OpenAI's Charter, a document we recently published that governs how we as an organization approach our mission and the broader community.

#1 We need to maximize AI's societal benefits and minimize its potential harms by building strong channels for dialogue between policy-makers and the community of researchers about ethical norms for responsible innovation in AI.

#2 We must support our academic institutions to allow us to meet the evident demands for AI talent and to ensure the US continues to define the frontier of AI research and applications, while also supporting the retraining of American workers to take advantage of this technological revolution.

#3 We must measure the progress and capabilities of AI to guide effective policy making.

#1: Ethical Norms.

Stakeholders in the development and deployment of AI technology, including the research community, actors in the private sector, and policy-makers, must jointly develop a set of ethical norms to govern their conduct and ensure that AI is both developed and deployed responsibly. If we deliberately pursue the creation of ethical norms around the research, development, and deployment of artificial intelligence, then we will be able to shape the culture in which the technology is developed, creating a shared sense of valid and invalid approaches to the technology, with specific best practices for specific areas. If we are successful in this then it will be significantly easier to apply regulations to artificial intelligence's consequences in the future as the community will have settled on a set of pre-agreed upon conventions that can become templates for law.

Without these norms it's likely that AI could be built or used in ways that cause harm to people or destabilize fundamental aspects of civil life. Last year, we saw the emergence of DeepFakes¹, technology based on artificial intelligence that made it relatively easy for people to take the face of a well-known individual, such as a celebrity, and superimpose it onto the faces of actors in pornographic films. This technology made it relatively easy for people to create unpleasant, unethical media, and while there is no indication anyone was harmed as a direct consequence of DeepFakes, I can assure you that the same technology could be used in other contexts - what if similar techniques were used to make it easy to take the face of a politician and superimpose it on another person in another context? We already have some examples of this occurring, such as a recent demonstration at MIT of President Trump's face being mapped onto the face of President Obama², and vice versa. The potential ramifications with regard to automated, synthetic propaganda are chilling and real.

At OpenAI, we have been thinking about these issues and, along with a wide set of stakeholders, recently published a report, the Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation³. In this report we have tried to think through some of the ways in which today's AI technologies could be re-purposed by malicious actors to cause harm, whether that be by using off-the-shelf AI technology to augment existing hacking techniques, or to use some of the aforementioned synthesis technologies to create more convincing 'advanced persistent threat'-style attacks, or to take soon-to-be-viable cleaning

¹ Oberoi, Gaurav. "Exploring DeepFakes." Hackernoon. March. 5, 2018.
<https://hackernoon.com/exploring-deepfakes-20c9947c22d9>

² Clark, Jack. A video recording of a demonstration of AI-based facial mapping. Twitter. March. 1, 2018.
<https://twitter.com/jackclarkSF/status/969285043502878722>

³ Brundage, Avin, et al. "The Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation." Arxiv. February. 20, 2018.
<https://arxiv.org/abs/1802.07228>

robots and convert them into bomb-delivery systems. Why did we do this? Because we wanted to highlight the 'dual use' nature of much of today's technology; the same AI technology that can be used to diagnose tumors from x-rays can also be used to train systems to surveil or target individuals, and because these capabilities are embodied in software, proliferate widely, run on standard compute hardware, and are developed by a global set of actors, then traditional control regimes and other policy tools don't seem to apply. Instead, we need to change the ethical norms with which developers approach their work on AI, so that they think twice before releasing something that can be trivially repurposed for negative uses, and so they try to use an 'adversary mindset' to view their own contributions through the perspective of a prospective attacker. I think this neatly illustrates how advances in AI capabilities also create new threats which are difficult to deal with through traditional regulatory tools, but may be dealt with by other means.

Another area for ethical norms, which OpenAI works on directly, is ensuring that the global development community is aware of and implements AI safety techniques in their own work to ensure that AI systems act in a reliable way in line with the intentions of their operators.

The government can and should, in our view, become more-active participants in the dialogue around developing shared ethical norms. As we've noted, the risks are quite substantial, so much so that we believe they're worthy of government attention, forecasting, and mitigation. Fostering the discussion around ethical norms is one of the ways that government can help reduce risk, and its impact-to-cost ratio is attractively high. The nature of threat forecasting and norm creation is by its very nature extremely interdisciplinary, and the government is well-positioned to bring together all of the stakeholders for productive discussions that lead to specific recommendations and clearly-identified focus areas. Furthermore, as AI becomes more central to information technology in general, the government will become a direct stakeholder in using it for some of the most sensitive applications possible---eg health care, law enforcement, and national security. This increases the importance of involving the government in these conversations and of government not only helping to define norms, but adopting them itself in its own practices.

Why do we think developing ethical norms will make a big difference? Because we know that it works. In the areas of bias, we have seen similar research in recent years in which people have highlighted how today's existing AI systems can exhibit biases, and the surfacing of these biases has usually led to substantive tweaks by the operators of the technology, as well as stimulating a valuable research discipline that provides a set of 'checks and balances' on AI development without the need for hard laws. We've also seen this happen in the area of algorithm deployment where organizations such as AI Now have embedded with communities likely to be affected by the use of automated algorithms and via joint research have devised recommendations that can be applied by the whole of the public sector. The more we can increase participation in these areas, the better prepared we will be for the significant challenges of this technology, and this will let us get ahead of some of its more obvious downsides so that we don't end up needing to do reactive regulation.

Ethical Norms: Recommendations:

My specific recommendations for next steps here are, beyond hearings like this, to have more dialogue between the AI community and political appointees and staffers here on the hill, as with specific groups within the government's agencies that are currently implementing their own AI plans and initiatives.

I would also be delighted to facilitate a workshop with people here in Washington to better understand the areas where people are particularly concerned about malicious actors re-purposing AI technology. Through such a workshop I would hope to better incorporate the views of the stakeholders represented here into the technical community's research into this important area, and thereby allow us to create norms that are sensitive to concerns you may be hearing from constituents or other parties.

#2: Strengthening the AI workforce.

Compared to 2013, there are now 4.5 times as many US jobs listed on the job search website 'Indeed.com' that require AI skills, according to the AI Index⁴. This is a symptom of a challenge the AI sector in America faces: demand for people with AI skills is outstripping the supply of people with those skills, which is causing the private sector to hire people away from academia and to hire students earlier in their careers, and we are not increasing funding to let academia keep pace with demand from students⁵; these three factors, combined together, potentially weaken AI education in the US in the future. The private sector has independently sought to meet this demand via the emergence of a range of online education courses for AI, and companies such as Google and Microsoft have also released tools and training courses to help other people transition into AI careers. These indicate a huge amount of demand for people with AI skills. Without aggressive investment into enabling basic scientific research, the US risks squandering its opportunity to lead in the development of AI and the education of its most important practitioners, and thereby might miss out on the upsides of the technology as well as the opportunity to shape both national and global regulations and norms for the technology.

AI will have a significant influence on the economy. The one trillion dollar question is whether it is going to be a good influence or a bad one for the workforce. Here, experts are divided. By 2020 the World Economic Forum believes AI may destroy around 7 million jobs worldwide while creating 2 million jobs, while Gartner believes it will destroy around 1,800,000 jobs while creating 2,300,000; by 2030 McKinsey believes it could destroy anywhere between 400 million

⁴ AI Index 2017 report, page 18. Accessed April. 10, 2018.
<http://aiindex.org/#report>

⁵ In 2011, the University of California at Berkeley received 474 applications for PhDs in artificial intelligence and was able to admit 13 students. By 2018 this had risen to 2229 applications with 46 admitted students, showing how demand has scaled faster than educational capacity.

and 800 million jobs and create anywhere between 555 and 890 million jobs.⁶ What these divergent estimates make clear is that AI will cause a significant amount of disruption in the labor markets around the world as some jobs are automated and other, new fields are created. But what we have learned from the current technology revolution is that many of the new jobs created require a significant set of skills than those which are automated. If we are to benefit as a nation and as a global community from AI then we must make sure this revolution is an inclusive one in which we equip as many people as possible with the skills needed to benefit from the changing economy, rather than be sidelined by it.

Strengthening the AI Workforce: Recommendations:

One of the best ways to support the basic research ecosystem - which all commercialization depends on - is to increase the number of funded PHD fellowships for AI across the country. This will make it easier for the United States to maintain its position as the global leader for AI education and will increase the chance of the world's next great artificial intelligence companies being founded here. We should pay particular attention to ensuring we fund students applying from abroad, as AI is of such opportunity it would be sensible to ensure that the US can educate and support the smartest people in the world. Government could choose to directly endow more funded PHD positions at universities that have displayed excellence in AI. This would have immediate beneficial effects and would avoid the additional overhead introduced by increasing grants which will generate significant logistical overhead on the part of professors that wish to apply for the new funding.

Government should take steps to support further commercialization and spin-outs of University research to further our already diverse ecosystem of AI startups, potentially via providing specialized funding to public universities to allow them to do this. As an example, the University of California at Berkeley recently formed The House⁷, a startup commercialization and funding entity which launched in 2016 and has subsequently backed more than 50 startups which have collectively raised over \$400 million. If government were able to supply relatively small amounts of capital to let public universities spin-up other, similar initiatives then it's fairly likely that the private sector would respond with additional funding and interest.

The government should invest more heavily in retraining programs both for its own workforce as well as for the workforce across America. We already know that in our current technology revolution the gains have been unequal, with certain regions, such as Silicon Valley, or the upper Pacific Northwest, or the North East of the United States, benefiting from job creation, and other regions languishing as jobs are automated and not replaced with new opportunities. This

⁶ Winick, Erin. "Every study we could find on what automation will do to jobs, in one chart". MIT Technology Review. January. 25, 2018.
<https://www.technologyreview.com/s/610005/every-study-we-could-find-on-what-automation-will-do-to-jobs-in-one-chart/>

⁷ The House. Accessed April. 16, 2018.
<https://thehouse.build/>

lack of equality drives discontent and reduces the faith of the worst-affected citizens in their political representatives which ultimately leads to a rise in extremism - we must avoid this.

#3 Measurement and AI Progress.

In my work on the AI Index it has become clear to me that we currently lack a bunch of the basic inputs we'd need to measure aspects of AI and how it is likely to evolve in the future. These measurement challenges include:

- A lack of knowledge about the rate of progress of self-driving cars as the companies developing them are keeping performance metrics secret as they view this as revealing proprietary data about their systems.
- A lack of standardized environments in which to test and evaluate robotics research and applications.
- Little information at both a state and country-level about the deployment of AI systems by both the government and the private sector.
- Little coordinated evaluation of the readiness of AI to be deployed in transformative areas like healthcare, and so on.

While I hope that the AI Index will motivate some further data gathering in these areas, I think the task is so large that it's an area where government could - and should - be more involved. If more people in government were focused on measuring and assessing the impact and progress of AI, then there would be more people in government aware of its progression and able to create smart policy to ensure it benefits as many people as possible.

Today, agencies like DARPA and IARPA are performing regular, subject-specific assessments of aspects of AI via hosting competitions and funding research, and NIST is developing some of the standards and assessment metrics as well. There are also bodies like the Congressional Research Service and the GAO which are looking at this area - last year I participated in a GAO-led report⁸ to try to assess AI's impacts and opportunities in a few areas across America.

Measurement and AI Progress: Recommendations:

The government should create and host more competitions to galvanize activity in the academic and private sectors, as DARPA did with its 'Cyber Grand Challenge' initiative in 2016, or DIUx did this year with its 'xView' dataset release. In particular, I can imagine that competitions relating to robot manufacturing, drone navigation and control, and the safety and predictability of AI systems would also serve to catalyze progress which would lead to commercial innovations and an increase in the robustness and usability of the technology. We should fund and encourage organizations, including those named above, to conduct more competitions in this area.

⁸ GAO, "Artificial Intelligence: Emerging Opportunities, Challenges, and Implications". March. 28, 2018. <https://www.gao.gov/products/GAO-18-142SP>

Additionally, we should be providing further funding for long-term AI measurement and analysis schemes, which could include empowering a specific agency to be the pre-eminent measurer of AI. For this, NIST may be a logical agency.

We may also want to explore even more unconventional ideas - for instance, might it be worth reviving the Office of Technology Assessment? OTA operated between 1972 and 1995 and during its lifetime produced more than 700 reports on technology areas of interest to the government. Having that kind of unbiased, bipartisan, science-driven organization would benefit the US today, as it would provide another means by which the government can educate itself about AI. Regardless, OTA provides a template for some of what can be achieved given a sufficiently powerful mandate and specific focus.

In particular, we should be seeking to measure and analyze the impact of AI on the economy and the workforce as, as discussed in section #2, this is an area with a diverse range of possible outcomes and where more knowledge will better prepare us as a nation to take advantage of this technology.

As well-meaning as I and my fellow panelists for this hearing are, I have a strong belief that there's no greater means of guaranteeing that you're getting the most correct view on a subject than by learning about it yourself. If the US government can take on a global leadership role in the measurement and forecasting of AI then it will be better positioned to identify the technologies challenges and opportunities, will be well positioned to create the competitions, datasets, and challenges that motivate progress in key areas, and can define global standards for what we, as a global community of actors, should be measuring and forecasting to ensure the continued stability of the world. There should be a concerted cross-agency effort to systematically analyze, measure, and forecast aspects of artificial intelligence deemed critical to the areas the agencies have purview over, and in areas that representatives hear frequent questions from constituents about.

Closing Statement:

If America can lead the world in the shaping of ethical norms around AI, strengthen its education system to ensure it remains the best place to study and develop AI technology, and ensure that the US government can be the most informed government in the world about AI's progress, then I think we'll be positioned to help not just American citizens, but the entire world benefit from this technology. The choice is, thankfully, ours to make.

Finally, I shall enclose the text of the OpenAI Charter, a document which we recently published that commits our own organization to a set of principles with which we'll approach the development of this technology. I think that another norm that we should all work on creating is making it easier and more acceptable for actors in the private sector to publicly commit themselves to standards that befit the immense impact of this technology.

OpenAI Charter:

OpenAI's mission is to ensure that artificial general intelligence (AGI) — by which we mean highly autonomous systems that outperform humans at most economically valuable work — benefits all of humanity. We will attempt to directly build safe and beneficial AGI, but will also consider our mission fulfilled if our work aids others to achieve this outcome. To that end, we commit to the following principles:

Broadly Distributed Benefits

- We commit to use any influence we obtain over AGI's deployment to ensure it is used for the benefit of all, and to avoid enabling uses of AI or AGI that harm humanity or unduly concentrate power.
- Our primary fiduciary duty is to humanity. We anticipate needing to marshal substantial resources to fulfill our mission, but will always diligently act to minimize conflicts of interest among our employees and stakeholders that could compromise broad benefit.

Long-Term Safety

- We are committed to doing the research required to make AGI safe, and to driving the broad adoption of such research across the AI community.
- We are concerned about late-stage AGI development becoming a competitive race without time for adequate safety precautions. Therefore, if a value-aligned, safety-conscious project comes close to building AGI before we do, we commit to stop competing with and start assisting this project. We will work out specifics in case-by-case agreements, but a typical triggering condition might be "a better-than-even chance of success in the next two years."

Technical Leadership

- To be effective at addressing AGI's impact on society, OpenAI must be on the cutting edge of AI capabilities — policy and safety advocacy alone would be insufficient.
- We believe that AI will have broad societal impact before AGI, and we'll strive to lead in those areas that are directly aligned with our mission and expertise.

Cooperative Orientation

- We will actively cooperate with other research and policy institutions; we seek to create a global community working together to address AGI's global challenges.
- We are committed to providing public goods that help society navigate the path to AGI. Today this includes publishing most of our AI research, but we expect that safety and security concerns will reduce our traditional publishing in the future, while increasing the importance of sharing safety, policy, and standards research.

Mr. HURD. Thank you, Mr. Clark. Well, you're in the right place. Measurement may not be flashy, but we talk about IT procurement as well, which isn't sexy either. So you're with good company.

Ms. Lyons, you're now recognized for 5 minutes.

STATEMENT OF TERAH LYONS

Ms. LYONS. Good afternoon. Chairman Hurd, Ranking Member Kelly, thank you for the opportunity to discuss a very important set of issues.

I am the executive director of the Partnership on Artificial Intelligence to Benefit People and Society, a 501(c)(3) nonprofit organization established to study and formulate best practices on AI technologies, to advance the public's understanding on AI, and to serve as an open platform for discussion and engagement about AI and its influences on people and society.

The Partnership is an unprecedented multistakeholder organization founded by some of the largest technology companies, in conjunction with a diverse set of cross-sector organizations spanning civil society and the not-for-profit community and academia. Since its establishment, the Partnership has grown to more than 50 partner organizations spanning three continents.

We believe that the formulation of the partnership could not have come at a more crucial time. As governments everywhere grapple with the implications of technology on citizens' rights and governance and as the research community increasingly emphasizes the need for multidisciplinary work focused on, not just the question of how we build technologies, but in some cases, whether to and also in what ways, the Partnership seeks to be a platform for collective reflection, and importantly, collective action.

My remarks this afternoon will focus, first, on some of the potential opportunities and challenges presented by artificial intelligence, and second, on how the Partnership hopes to engage with policymakers with industry, the research community, and other stakeholders. Artificial intelligence technologies present a significant opportunity for the United States and for the world to address some of humanity's most pressing and large-scale challenges, to generate economic growth and prosperity, and to raise the quality of human life everywhere.

While the promise of AI applied to some domains is still distant, AI is already being used to solve important challenges. In healthcare, already mentioned, AI systems are increasingly able to recognize patterns in the medical field helping human experts interpret and scan and detect cancers. These methods will only become more effective as large datasets become more widely available. And beyond healthcare, AI has important applications in environmental conservation, education, economic inclusion, accessibility, and mobility, among other areas.

As AI continues to develop, researchers and practitioners must ensure that AI-enabled systems are safe, that they can work effectively with people and benefit all parts of society, and that their operation will remain consistent and aligned with human values and aspirations. World-changing technologies need to be applied and ushered in with corresponding social responsibility, including attention paid to the impacts that it has on people's lives.

For example, as technologies are applied in areas like criminal justice, it is critical for the Partnership to raise and address concerns related to the inevitable bias and datasets used to train algorithms. It's also critical for us to engage with those using such algorithms in the justice system so that they understand the limits of these technologies and how they work.

Good intentions too are not enough to ensure positive outcomes. We need to ensure that ethics are put into practice when AI technologies are applied in the real world and that they reflect the priorities and needs of the communities that they serve. This won't happen by accident. It requires a commitment from developers and other stakeholders who create and influence technology to engage with broader society, working together to predict and direct AI's benefits and to mitigate potential harms. Identifying and taking action on high-priority questions for AI research, development, and governance will require the diverse perspectives and resources of a range of different stakeholders, both inside and outside of the partnership on AI community.

There are several ways in which we are delivering this. A key aspect of this work of the Partnership has so far taken the form of a series of working groups which we have established to approach three of our six thematic pillars, with the other three to follow soon. These first working groups are on safety critical artificial intelligence; fairness, transparency, and accountability in AI; and also AI labor in the economy.

The Partnership will also tackle questions that we think need to be addressed urgently in the field and are ripe for collective action by a group of interests and expertise as widespread and diverse as ours. Our work will take different forms and could include research, standards development, policy recommendations, best practice guidelines, or codes of conduct. Most of all, we hope to provide policymakers and the general public with the information they need to be agile, adaptive, and aware of technology developments so that they can hold technologists accountable for upholding ethical standards in research and development and better understand how these technologies affect them.

We are encouraged by these hearings and the interest in policymakers in the U.S. and worldwide both toward understanding the current state of AI and the future impacts it may have.

I thank you for your time, and I look forward to questions.

[Prepared statement of Ms. Lyons follows:]

Written Testimony of Terah Lyons
Executive Director
The Partnership on Artificial Intelligence to Benefit People & Society
House of Representatives Oversight & Government Reform Committee
Subcommittee on Information Technology
Hearing on
Game Changers: Artificial Intelligence Part III – AI and Public Policy

April 18, 2018

Chairman Hurd, Ranking Member Kelly, and Members of the Subcommittee, thank you for the opportunity to discuss a unique and unprecedented multi-stakeholder body—the Partnership on Artificial Intelligence to Benefit People & Society (PAI, or the Partnership)—and to provide PAI’s perspective on the important issues raised by artificial intelligence (AI). We appreciate the Subcommittee’s interest in the potential challenges and opportunities presented by artificial intelligence, the possible barriers to effective development and deployment of artificial intelligence, and the possible solutions in response to these challenges and barriers. The Partnership tackles both substantive AI and AI governance challenges in its mission to resolve such questions.

Like you, we believe that AI is and will be beneficial to people and society—providing valuable services, improving efficiency for citizens, companies, and governments, helping people with a wide range of individual goals, and even saving lives. We also recognize that AI invites important questions, such as the potential impact on our workforce; the possibility of entrenching bias; and the ramifications for privacy. One thing is clear—if we allow a gulf to open up between the entities developing AI technologies and the people and societies that are impacted by them, we will both limit progress and cause harm in the process. Every organization involved in the Partnership on AI has accepted that premise as part of their membership, and each organization is committed to maximizing the benefits of AI while working to address its potential challenges.

In fact, our early work spans several of your priority issues. PAI’s initial Working Groups—the main bodies through which PAI examines AI best practices, marshals resources toward high-priority AI research, and identifies aspirational pro-social AI projects—are focused on the labor impacts of automation; fairness, accountability, and transparency of AI systems; and the ethical and effective application of AI in safety-critical scenarios.

You are also interested in exploring how government and other agents can most effectively address those challenges. We believe that PAI’s multi-stakeholder approach is vital in delivering positive outcomes in technology development and technology governance. We must

work together to find a path forward in understanding and delivering solutions to the complex challenges associated with AI development—and in maximizing the benefits of AI—by leveraging diverse interests, expertise, and perspectives. Congress’ support for similarly diverse and nuanced conversations will be important to this work, now, and in the future.

We would like to be clear that PAI is not a lobbying organization. PAI does, however, intend to be a *resource* to policymakers—for instance, in conducting research that informs AI best practices and exploring the societal consequences of certain AI systems, as well as policies around the development and usage of those systems. Toward those ends, PAI is fortunate to have a diverse membership of experts which spans sectors and fields – AI researchers and engineers, companies, academics from a wide range of disciplines, advocates for interests related to human rights, civil liberties, and other leaders from civil society. It is a deep and growing bench, with unique perspectives and an emphasis on equity informed by its equal representation of for-profit and not-for-profit members on our Board of Directors. We believe that our diversity is vital to the success, sustainability, and durability of the organization and its mission, and critical to effective and inclusive multi-stakeholder engagement.

My testimony today will discuss:

- **The promise of AI**
- **The formation of PAI and the need it serves**
- **PAI’s four goals**
- **PAI’s early work**
- **What is next for PAI and AI governance**

The Promise of AI

Artificial intelligence technologies present a significant opportunity for the United States and for the world to address some of humanity’s most pressing and large-scale challenges, to generate economic growth and prosperity, and to raise the quality of human life everywhere.

While we might be years or decades away from fully realizing that opportunity, it is important to recognize that AI technologies are already solving important challenges today. Researchers are building systems that show increasing abilities to conduct pattern recognition in healthcare, including augmenting human experts in interpreting radiologic studies and detecting cancers from images at human expert levels, with the increasing availability of large datasets from which algorithms can learn. AI advances also show promise at providing new efficiencies and optimizations, in applications ranging from the products we purchase and the transportation technology and infrastructure we use, to the ways in which we consume energy. AI technologies are also leading to enhanced understanding of global systems at scale through the collection and analysis of data about global health, the environment, and more.

There are many similar examples about how AI is being used in beneficial ways today with the promise of even more advancements in the future. To illustrate the trajectory of some of these innovations, AI has the potential to be a new economic engine to drive growth, opportunity, and prosperity. Over the last decade, a global explosion of smartphones have

become the infrastructure for a new app economy, catalyzing new businesses, both big and small. The next decade could see the creation of an AI economy, ushered in by a new wave of efficiency and optimization. Healthcare, nutrition, environmental conservation, economic inclusion, and accessibility and mobility, among other domains, all have the potential to see radical shifts for the positive as a result of AI.

In order to ensure a strong future where we harness the latent value of AI methods, we must begin thinking carefully, starting now, about the challenges of development and implementation posed by these technologies and how we can address them to ensure a future that benefits us all.

The Formation of PAI and the Need it Serves

With the promises and risks of AI in mind, PAI was established to study and formulate best practices on AI technologies, to advance the public's understanding of AI, and to serve as an open platform for discussion and engagement about AI and its influences on people and society.

We are at a critical juncture for AI's development and its applications. The development of AI technologies is now intersecting with the application of these systems and techniques. Advancements in algorithms, computing power, and rich data, are helping the field to solve technical challenges that will improve domains such as perception, natural language understanding, and robotics, bringing great value in the years ahead.

However, with technological advancement comes concerns and challenges associated with responsible development and the impacts of technologies on people's lives. These concerns include the safety of AI technologies, the fairness and transparency of systems, and the intentional as well as inadvertent influences of AI on people and society, including potential influences on privacy, democracy, criminal justice, and human rights. We must also be mindful of concerns associated with the potential unintended harms that "AI-for-good" applied in some circumstances and manners may impose, and must work to ensure that good intentions in technology application and deployment translate into positive outcomes for the people and communities they impact.

The Partnership's formation was first rooted in leading technologists collectively understanding the need for AI best practices, dialogue, and common understanding related to issues such as these. From the outset, the Partnership was designed as a multistakeholder entity, and from its founding has grown to include a community of technology companies, academic institutions, civil society organizations and other types of not-for-profit institutions, and increasingly global perspectives. Diversity of thought is critical to PAI's success, as is considering the impact of AI on diverse constituencies. An organization comprised of for-profit and not-for-profit members has more nuanced perspective to draw upon in identifying and effectively addressing these challenges. To that end, we aspire to an even more representative community than PAI currently reflects, with engagement from even more diverse interests than currently represented, specific application areas and sectors not currently involved in our Partner

community, and underrepresented interests that may give voice to constituencies needed but not often heard in matters of technology governance.

The Partnership's goal is to involve deep subject-matter experts, across a variety of disciplines, in our work. PAI also aims to focus both on identifying and developing issues and topic areas ripe for consensus, as well as grappling with more challenging and nuanced questions associated with the development and deployment of AI. Through this issue identification, and through enumerating, interrogating, and leveraging the varying perspectives and interests of our multi-stakeholder body, it is the Partnership's aspiration to incite critical reflection about the impacts of AI on people and society, and to constructively shape practices, norms, expectations, and behavior related to responsible AI development.

The Partnership's activities are deliberately determined by its Partners, but from the outset of the organization, the hope and intention has been to create a place for open critique and reflection. Crucially, the Partnership is an independent organization; though supported and shaped by our Partner community, the Partnership is ultimately more than the sum of its parts, and will make independent determinations to which its Partners will collectively contribute, but never individually dictate.

It is our intention that PAI will be collaborative, constructive, and work openly with all, both those who find themselves aligning with recommendations and guidance on best practices and those who differ in assessments—or who are skeptical of our work. Crucially, PAI has been explicitly designed to bring together varying perspectives in a structure that ensures balanced governance by diverse stakeholders. In the same manner that PAI welcomes and encourages thoughtful debate among its members, PAI also welcomes debate among society on the best ways to engage these topics.

PAI's Goals

The formation of PAI was based on a clear set of goals and overarching Tenets (see Appendix). Based on these goals and founding principles, substantive Thematic Pillars were developed to guide the organization's work, in pursuit of the following outcomes:

- ***Develop and Share Best Practices***

Support research, discussions, identification, sharing, and recommendation of best practices in the research, development, testing, and fielding of AI technologies. Address such areas as fairness and inclusivity, explanation and transparency, security and privacy, values and ethics, collaboration between people and AI systems, interoperability of systems, and of the trustworthiness, reliability, containment, safety, and robustness of the technology.

- ***Provide an Open and Inclusive Platform for Discussion and Engagement***

Create and support opportunities for AI researchers and key stakeholders, including people in technology, law, policy, government, civil liberties, and the greater public, to communicate directly and openly with each other about relevant issues to AI and its influences

on people and society. Ensure that key stakeholders have the knowledge, resources, and overall capacity to participate fully.

- ***Advance Public Understanding***

Advance public understanding and awareness of AI by multiple constituencies, including writing and other communications on core technologies, potential benefits, and costs. Act as a trusted and expert point of contact as questions, concerns, and aspirations arise from the public and others in the area of AI. Regularly update key constituents on the current state of AI progress.

- ***Identify and Foster Aspirational Efforts in AI for Socially Beneficial Purposes***

Seek out, support, celebrate, and highlight aspirational efforts in AI for socially benevolent applications. Identify areas of untapped opportunity, including promising technologies and applications not being explored by academia and industry R&D.

We believe that artificial intelligence technologies hold great promise for raising the quality of people's lives and can be leveraged to help humanity address important global challenges such as climate change, food, inequality, health, and education. Our mission is in service of unlocking this potential through intentional, inclusive, multidisciplinary analysis, design, and accountability.

PAI's Early Work

PAI's efforts and early-stage programs are currently organized around a set of Thematic Pillars reflecting key challenges and opportunities in AI:

- ***AI Labor, and the Economy***

AI advances will undoubtedly have multiple influences on the distribution of jobs and nature of work. While advances promise to inject great value into the economy, they can also be the source of disruptions as new kinds of work are created and other types of work become less needed due to automation.

Discussions are rising on the best approaches to minimizing potential disruptions, making sure that the fruits of AI advances are widely shared and competition and innovation is encouraged and not stifled.

- ***Safety-Critical AI***

Advances in AI have the potential to improve outcomes, enhance quality, and reduce costs in such safety-critical areas as healthcare and transportation. Effective and careful applications of pattern recognition, automated decision making, and robotic systems show promise for enhancing the quality of life and preventing thousands of needless deaths.

However, where AI tools are used to supplement or replace human decision-making, we must be sure that they are safe, trustworthy, and aligned with the ethics and preferences of people who are influenced by their actions.

- ***Fair, Transparent, and Accountable AI***

AI has the potential to provide societal value by recognizing patterns and drawing inferences from rich amounts of data. Data can be harnessed to develop useful diagnostic systems and recommendation engines, and to support people in making breakthroughs in such areas as biomedicine, public health, safety, criminal justice, education, and sustainability.

While such results promise to provide great value, we need to be sensitive to the hidden assumptions and biases in data as well as the biases reflected in decisions about system design. This can lead to actions and recommendations that replicate those biases, and suffer from serious blind spots. Such failings can lead to the unfair and systematic exclusion of groups of people from consequential resources and services, such as in financial services, job opportunities, and education.

Researchers, officials, and the public should be sensitive to these possibilities and we should seek to develop methods that detect and correct those errors and biases, not replicate them. We also need to work to develop systems that can explain the rationale for inferences.

- ***Collaborations Between Humans and AI Systems***

A promising area of AI is the design of systems that augment the perception, cognition, and problem-solving abilities of people. Examples include the use of AI technologies to help physicians make more timely and accurate diagnoses and assistance provided to drivers of cars to help them to avoid dangerous situations and crashes.

Opportunities for R&D and for the development of best practices on AI-human collaboration include methods that provide people with clarity about the understandings and confidence that AI systems have about situations, means for coordinating human and AI contributions to problem solving, and enabling AI systems to work with people to resolve uncertainties about human goals.

- ***Social and Societal Influences of AI***

AI advances will touch people and society in numerous ways, including potential influences on privacy, democracy, criminal justice, and human rights. For example, while technologies that personalize information and that assist people with recommendations can provide people with valuable assistance, they could also inadvertently or deliberately manipulate people and influence opinions.

We seek to promote thoughtful collaboration and open dialogue about the potential subtle and salient influences of AI on people and society.

- ***AI and Social Good***

AI offers great potential for promoting the public good, for example in the realms of education, housing, public health, and sustainability. We see great value in collaborating with public and private organizations, including academia, scientific societies, NGOs, social entrepreneurs, and interested private citizens to promote discussions and catalyze efforts to address society's most pressing challenges.

Some of these projects may address deep societal challenges and will be moonshots – ambitious big bets that could have far-reaching impacts. Others may be creative ideas that could quickly produce positive results by harnessing AI advances. Beyond the specified thematic pillars, we also seek to convene and support projects that resonate with the tenets of our organization. We are particularly interested in supporting people and organizations that can benefit from the Partnership’s diverse range of stakeholders, and are welcoming input from stakeholders regarding what work would be most valuable for an entity like PAI to undertake. We are open-minded about the forms that these efforts will take.

Working Groups

The organization’s earliest programming work has established Working Groups associated with three of the above Thematic Pillar areas: AI, Labor, and the Economy; Safety-Critical AI; and Fair, Transparent, and Accountable AI. Through these structures, Partners opt into Working Groups addressing each Pillar. Co-Chairs—one each representing a for-profit and from a not-for-profit Partner—lead each Working Group. This diverse leadership helps to ensure that the Partnership empowers and draws upon the full breadth of its perspectives, and is designed to help ensure that all members have a seat at the table. Over time, we will launch additional Working Groups focused on other areas of work. And as the work of these sub-communities develops, we look forward to updating the public as to PAI’s organizational priorities, projects, and work products.

Providing an Open and Inclusive Platform for Engagement

From the founding of the Partnership on AI, we have been optimistic about the high potential of our Working Groups to generate meaningful insights and to incite cross-sector, multidisciplinary reflection about opportunities and challenges in the AI field. As an organization, the Partnership believes that transparency, diversity, and inclusivity are essential in advancing this work. From the outset, we intend to create a culture where Working Group members listen closely to those with whom they *disagree*—understanding that a central value of the Partnership is in our capacity to learn from differing interests, perspectives, and areas of expertise. Together, we will work to leverage the power of our diversity to create an environment which promotes meaningful collective reflection, and produces impactful decisions and outputs. The project of maintaining the strength and integrity of our community will rely on commitment to these ideals from across our organization.

The Partnership is dedicated to creating and supporting opportunities for AI researchers and key stakeholders, including people working in technology, law, policy, government, advocacy—and the broader public—to communicate directly and openly with each other about relevant issues to AI and its influences on people and society. The Partnership is also dedicated to ensuring that its stakeholders have the knowledge, resources, and overall capacity to participate fully. To ensure that valued perspectives are never marginalized, the Partnership is in the midst of developing a full program of support and capacity-building for civil society organizations (CSOs), which face unique challenges in supporting the type of work, often voluntary, required of meaningful participation in a multistakeholder organization like the Partnership. Though the organization is in early stages of developing this program area, the

Partnership is committed to fostering a sustainable multistakeholder model which enables the full participation of interested Partners.

As citizens, communities, and businesses across the U.S. and around the world grapple with important challenges associated with the development of AI technologies, we would invite members of Congress and the broader public to engage with PAI to ensure we are focusing our attention and the attention of our Partners, accordingly.

What is Next for PAI and AI Governance

The work of the Partnership is motivated by the belief that because the societal implications of AI development and deployment are complex and multifaceted we cannot determine solutions alone—neither as singular entities, or as a singular discipline. As the AI field and multidisciplinary understandings of its impact evolve over time, the Partnership will be a place where stakeholders remain engaged on these topics, and work to update, refine, and advance perspectives and solutions to associated challenges.

The Partnership will be offering perspectives on topics like those related to our Thematic Pillars and current and future Working Groups. We are committed to supporting the following immediate areas of work:

- **Engagement of Experts**

The regular engagement of experts across multiple disciplines (including but not limited to psychology, philosophy, economics, finance, sociology, public policy, and law) to discuss and provide guidance on emerging issues related to the impact of AI on society.

- **Engagement of Other Stakeholders**

The engagement of AI users and developers, as well as representatives of industry sectors that may be impacted by AI (such as healthcare, financial services, transportation, commerce, manufacturing, telecommunications, and media) to support best practices in the research, development, and use of AI technology within specific domains.

- **Third-Party Support**

The design, execution, and financial support of objective third-party studies on best practices for the ethics, safety, fairness, inclusiveness, trust, and robustness of AI research, applications, and services. The identification and celebration of important work in these fields. The support of aspirational projects in AI that would greatly benefit people and society.

- **Informational Material Development**

The development of informational materials on the current and future likely trajectories of research and development in core AI and related disciplines.

Conclusion

We are excited about the future, and recognize that the Partnership is an ambitious project with high expectations both from its members and the greater public. With the launch of

our first Working Groups, we are getting to the pressing work at hand and look forward to making a valuable and lasting contribution to advancing trusted, ethical, inclusive, multidisciplinary, and beneficial AI.

Thank you for this opportunity to speak about the Partnership on AI, its goals, and its efforts. We look forward to being a resource for Congress, other global policymaking bodies, and the public, to help ensure that the benefits of AI are realized and the challenges are addressed as PAI pursues its mission.

ADDITIONAL MATERIALS

Tenets of the Partnership on AI

Members of the Partnership on AI share the following tenets:

1. We will seek to ensure that AI technologies benefit and empower as many people as possible.
2. We will educate and listen to the public and actively engage stakeholders to seek their feedback on our focus, inform them of our work, and address their questions.
3. We are committed to open research and dialogue on the ethical, social, economic, and legal implications of AI.
4. We believe that AI research and development efforts need to be actively engaged with and accountable to a broad range of stakeholders.
5. We will engage with and have representation from stakeholders in the business community to help ensure that domain-specific concerns and opportunities are understood and addressed.
6. We will work to maximize the benefits and address the potential challenges of AI technologies, by:
 - a. Working to protect the privacy and security of individuals.
 - b. Striving to understand and respect the interests of all parties that may be impacted by AI advances.
 - c. Working to ensure that AI research and engineering communities remain socially responsible, sensitive, and engaged directly with the potential influences of AI technologies on wider society.
 - d. Ensuring that AI research and technology is robust, reliable, trustworthy, and operates within secure constraints.
 - e. Opposing development and use of AI technologies that would violate international conventions or human rights, and promoting safeguards and technologies that do no harm.
7. We believe that it is important for the operation of AI systems to be understandable and interpretable by people, for purposes of explaining the technology.

8. We strive to create a culture of cooperation, trust, and openness among AI scientists and engineers to help us all better achieve these goals.

Members of the Partnership on AI

Trust is critical to the adoption of AI and realization of its benefits. Developing a framework for AI best practices – for it to work, for public acceptance, to make sure we have done it right – must include relevant and impacted stakeholders. A paramount directive of PAI is that its membership be reflective, in an equitable way, of both profit and non-profit entities. We also believe that, given the global nature of technology, that our membership include organizations from around the world.

Since we came together to form PAI we have been uplifted by support from around the world, with enthusiasm about our mission, tenets, and goals coming from companies, nonprofits, and academics alike. We continue to work to build a diverse, multi-stakeholder organization for open and constructive dialogue on AI.

At present, PAI has more than 50 organizations as members, with more than 120 representatives of these organizations contributing to its first Working Groups. Membership comprises civil society organizations, advocacy organizations, academic research laboratories and institutes, and for-profit organizations, including six of some of the largest technology companies in the world. The Board represents similarly diverse perspectives, with members hailing from organizations including large companies, as well as representatives of the ACLU, the MacArthur Foundation, OpenAI, and The Association for the Advancement of Artificial Intelligence. We think there is strength in this early ideological diversity and look forward to continuing to build upon it with representatives from new industries, new geographies, and new perspectives not often heard in discussions related to technology governance.

A complete list of Partners and Board members can be found at:

<https://www.partnershiponai.org/partners/>

<https://www.partnershiponai.org/board-of-directors/>

Mr. HURD. I appreciate you, Ms. Lyons.

Dr. Buchanan, you're now recognized for 5 minutes for your opening remarks.

STATEMENT OF BEN BUCHANAN

Mr. BUCHANAN. Thank you, Chairman Hurd and Ranking Member Kelly, for holding this important hearing and for inviting me to testify. As you mentioned, I'm a fellow at Harvard University's Belfer Center for Science and International Affairs, and my research focus is on how nations deploy technology, in particular, cybersecurity, including offensive cyber capabilities and artificial intelligence.

Recently, with my friend and colleague, Taylor Miller, of the Icahn School of Medicine at Mount Sinai, we published a report entitled, "Machine Learning for Policymakers." And to help open today's hearing, I would like to make three points: one on privacy, one on cybersecurity, and one on economic impact. And I'll try to tailor this to not be repetitive. I think we're in agreement on a lot of these areas.

To simplify a little bit, we can think about modern artificial intelligence as relying on a triad of parts: some data, some computing power, and some machine learning algorithms. And while we've seen remarkable advances on the computing and learning algorithm side, I think for policymakers such as yourselves, it's data that's most important to understand. And data is the fuel of machine learning systems. Without this data, the systems sometimes produce results that are embarrassingly wrong and incorrect.

Gathering relevant and representative data for training, development, and testing purposes is a key part of building modern artificial intelligence technology. On balance, the more data that is fed into a machine learning system, the more effective it will be. It is no exaggeration to say that there are probably many economic, scientific, and technological breakthroughs that have not yet occurred because we have not assembled the right data sources and right datasets.

However, there is a catch and a substantial one. Much of that data that might, and I emphasize might, be useful for future machine learning systems is intensely personal, revealing, and appropriately private. Too frequently, the allure of gathering more data to feed a machine learning system distracts from the harms that collecting that data brings. There is a risk of breaches by hackers, of misuse by those who collect or store the data, and of secondary use in which data is collected for one purpose and later reappropriated for another.

Frequently, attempts at anonymization do not work nearly as well as promised. It suffices to say that, in my view, any company or government agency collecting large amounts of data is assuming an enormous responsibility. Too often, these collectors fall far short of meeting that responsibility. And yet, in an era of increased artificial intelligence, the incentive to collect ever more data is only going to grow.

And technology cannot replace policy, but some important technological innovations can offer mitigation to this problem. Technology such as differential privacy; that approach can ensure that

large datasets retain a great deal of their value, but protecting the privacy of any one individual member. On-device processing can reduce the aggregation of data in the first place. This is an area in which much remains to be done.

Second, AI is poised to make a significant impact in cybersecurity, potentially redefining key parts of the entire industry. Automation on offense and on defense is an area of enormous significance. We already heard about the DARPA grand cyber challenge, which I agree was a significant, seminal event, and we've certainly seen what I would describe as the beginnings of significant automations of cyber attacks in the wild.

In the long run, it's uncertain whether increased automation will give a decisive cybersecurity advantage to hackers or to network defenders, but there is no doubt of its immediate and growing relevance.

AI systems also pose new kinds of cybersecurity challenges. Most significant among these is the field of adversarial learning in which the learning systems themselves can be manipulated oftentimes by what we call poisoned datasets to produce results that are inaccurate and sometimes very dangerous. And that's another area which is very nascent and not nearly as developed as mainstream cybersecurity literature. Again, much more remains to be done.

A more general concern is AI safety. And this conjures up notions of Terminator and AI systems that will take over the world. In practice, it is often far more nuanced and far more subtle than that, though the risk is still quite severe. I think it is fair to say that we have barely scratched the surface of important safety and basic security research that can be done in AI, and this is an area, as my fellow witnesses suggest, in which the United States should be a leader.

Third, AI will have significant economic effects. My colleagues here have discussed many of them already. The ranking member mentioned two notable studies. I would point you to two other studies, both I believe by MIT economists, which show that while theory often predicts a job's loss will be quickly replaced, in practice, at least in that one instance, that did not immediately occur.

With that, I will leave it there, and I look forward to your questions. Thank you.

[Prepared statement of Mr. Buchanan follows:]

Game Changers: Artificial Intelligence Part III

House Oversight Committee, Subcommittee on IT

Prepared Testimony and Statement for the Record of Ben Buchanan

Postdoctoral Fellow, Belfer Center Cybersecurity Project
Harvard University

Thank you, Chairman Hurd and Ranking Member Kelly, for holding this important hearing and for inviting me to testify.

My name is Ben Buchanan. I am a fellow at Harvard University's Belfer Center for Science and International Affairs, and a Global Fellow at the Woodrow Wilson International Center for Scholars. My research specialty is examining how nations deploy technology, and I especially focus on cybersecurity and artificial intelligence. Recently, with Taylor Miller of the Icahn School of Medicine at Mount Sinai, I co-authored a paper entitled "Machine Learning for Policymakers."¹

To help open today's hearing on artificial intelligence, I'd like to make three points: one on privacy, one on cybersecurity, and one on economic impact.

First, to simplify a bit, we can think of most modern artificial intelligence systems as relying on a triad of pillars: data, computing power, and learning algorithms. While we have seen remarkable advances in computer hardware and machine learning software, for many policy purposes it is the role of data that is most vital to understand. Data is the fuel of machine learning systems; without it, these systems produce embarrassingly poor results. Gathering relevant and representative data for training, development, and testing purposes is a key part of building modern artificial intelligence technology. On balance, the more data that is fed into a machine learning system, the more effective it will be. It is no exaggeration to say that there are probably many economic, scientific, and technological breakthroughs that have not yet occurred because the right data sets have not yet been assembled.

However, there is a catch, and a substantial one: much of the data that might—and I emphasize *might*—be useful for future machine learning systems is intensely personal, revealing, and appropriately private. Too frequently, the allure of gathering more data to feed a machine learning system distracts from the harms that collecting that data brings. There is the risk of breaches by hackers, of misuse by those who collect it or access it, and of secondary use—in which data collected for one purpose is later re-appropriated for another. Frequently, attempts at anonymization do not work nearly as well as promised. It suffices to say that, in my view, any company or government agency collecting large amounts of data is assuming an enormous

¹ Buchanan, Ben and Taylor Miller. "Machine Learning for Policymakers." *Belfer Center for Science and International Affairs* (2017),

<https://www.belfercenter.org/sites/default/files/files/publication/MachineLearningforPolicymakers.pdf>.

² Michele Banko and Eric Brill, "Scaling to Very Very Large Corpora for Natural Language Disambiguation" (paper presented at 'Proceedings of the 39th Annual Meeting on Association for Computational Linguistics', 2001).

responsibility. Too often, these collectors fall far short of meeting that responsibility. And yet, in an era of increased artificial intelligence, the incentive to collect ever more data is only going to grow.

Technology cannot replace policy, but some important innovations offer some mitigation to this problem. Technical approaches such as differential privacy ensure that the particular data of any one individual is obscured but that large data sets retain almost all of their value. On-device processing reduces the amount of data transmitted back to central servers and makes interception and aggregation of private information harder. But these technological advances are too infrequently deployed, especially when they conflict with short-term financial interests. This is an area in which much remains to be done.

Second, AI is poised to make a significant impact in cybersecurity, potentially redefining key parts of the industry. Automation on offense and defense is an area of enormous significance. The most high-profile example of this is the DARPA Grand Cyber Challenge, performed live at the DEF CON hacking conference in 2016, in which automated computer systems played both offense and defense against one another in a hacking competition. In the long run, it is uncertain whether increased automation will give a decisive cybersecurity advantage to hackers or to defenders, but there is no doubt of its immediate relevance.³

AI systems also pose new kinds of cybersecurity challenges. Most significant among these is the field known as adversarial learning, in which the learning mechanisms of algorithms can be misled. This can cause AI systems to make bizarre and unpredictable decisions. The more powerful the system, the potent such a mistake can be. Cybersecurity is challenging enough to begin with; adding in modern AI technology such as deep learning only makes it more so. Research in the world of AI system security is fairly early-stage, especially compared to the much more developed body of mainline cybersecurity research and best practices.

A more general security concern is AI safety. AI safety is the field of research and development that ensures that AI systems, once deployed, remained aligned with the original interests of their designers and do not pose unanticipated threats. This is not a question of Terminator scenarios. Rather it is usually far subtler, but vitally important and too frequently neglected.⁴ I think it is fair to say we have barely scratched the surface of the important safety and basic security research that can be done in AI, and that the United States should be a leader in these areas.

Third, AI will have significant economic effects. Some of these, to be sure, will be positive. On the other hand, some will be quite negative. The drumbeat of beneficial technological progress has always brought some upheaval with it. Nonetheless, my concern in this instance is that the further development and integration of AI will occur at such a rapid rate that its economic impacts might be hard to anticipate and counteract where appropriate. Some research already clearly supports this view: one major Oxford study cited in a landmark White House report suggests that 47% of jobs

³ Schneier, Bruce. "Artificial Intelligence and the Attack/Defense Balance." *IEEE Security and Privacy* (2018) https://www.schneier.com/essays/archives/2018/03/artificial_intelligence.html.

⁴ For one of the leading works in this area, see Amodei, Dario, et. al. "Concrete problems in AI safety." *arXiv preprint:1606.06565* (2016). For further development of these ideas and others, see "Example Topics: AI Alignment." *Open Philanthropy* (2017). <https://www.openphilanthropy.org/focus/global-catastrophic-risks/potential-risks-advanced-artificial-intelligence/open-philanthropy-project-ai-fellows-program#examples>.

could be threatened by machines, while another McKinsey analysis places the number at 30%.⁵ Other studies bear out a worrying trend: even when economic theory predict jobs lost to automation would be replaced, later empirical analysis suggests that this replacement did not occur in practice.⁶

We must make sure that our current and future workforce is competitive in an age in which AI will be ubiquitous, when it is as broadly integrated throughout our society as electricity, according to Stanford professor and leading AI researcher Andrew Ng. We should recognize that many American students and workers will need to have an understanding of computer science, statistics, and machine learning in order to be most competitive in the global economy; too often these subjects are not taught in our schools, or not taught well and with appropriate resources. Innovations like Massive Open Online Courses, or MOOCs, have been immensely important in helping Americans acquire these valuable skills, but government and formal educational settings have important roles to play as well. Other nations, such as China, have begun to dramatically invest in these areas of education and research, and we must do the same.

Simply put, AI is exciting, economically vital, and geopolitically important. Maximizing its potential will require a whole of society effort. I appreciate your efforts in holding this series of hearings, and I look forward to your questions.

⁵ Frey, Carl and Michael Osborne, "The Future of Employment: How Susceptible Are Jobs to Computerization," Oxford University (2013) (http://www.oxfordmartin.ox.ac.uk/downloads/academic/The_Future_of_Employment.pdf). See also, Manyika, James, et. al. "Jobs Lost, Jobs Gained: Workforce Transitions in a Time of Automation." *McKinsey Global Institute* (2017), <https://www.mckinsey.com/~media/McKinsey/Global%20Themes/Future%20of%20Organizations/What%20the%20future%20of%20work%20will%20mean%20for%20jobs%20skills%20and%20wages/MGI-Jobs-Lost-Jobs-Gained-Report-December-6-2017.ashx>.

⁶ Acemoglu, Daron, and Pascual Restrepo. "The Race between Machine and Man: Implications of Technology for Growth, Factor Shares and Employment", *NBER* (2016), <https://www.nber.org/papers/w22252.pdf>. Acemoglu, Daron, and Pascual Restrepo. "Robots and Jobs: Evidence from US Labor Markets." *NBER* (2017) <https://www.nber.org/papers/w23285>.

Mr. HURD. Thank you, Dr. Buchanan.

And I'll recognize the ranking member for 5 minutes or so for your first round of questions.

Ms. KELLY. Thank you, Mr. Chairman, and thank you to the witnesses again.

The recent news that Cambridge Analytica had improperly obtained the personal data of up to 87 million Facebook users highlights the challenges to privacy when companies collect large amounts of personal information for use in AI systems.

Dr. Buchanan, in your written testimony, you state, and I quote, that "much of the data that might—and I emphasize might—be useful for future machine learning systems is intensely personal, revealing, and appropriately private."

Is that right? You just said that.

Mr. BUCHANAN. That's correct, Congresswoman.

Ms. KELLY. And can you explain for us what types of risks and threats consumers are exposed to when their personal information is collected and used in AI systems?

Mr. BUCHANAN. Sure. As you'd expect, Congresswoman, it would depend on the data. Certainly, some financial data, if it were to be part of a breach, would lead to potential identity theft. There's also data revealed in terms of preferences and interests that many members of society might want to keep appropriately private. We've heard a lot about AI in medical systems. Many people want to keep their medical data private.

So I think it depends on the data, but there's no doubt that, in my view, if a company or government organization cannot protect the data, it should not collect the data.

Ms. KELLY. Okay. In light of these risks and your assessment on the majority of companies that do collect and use personal data for their AI systems, are they taking adequate steps to protect the privacy of citizens?

Mr. BUCHANAN. Speaking as a generalization, I think we have a long way to go. Certainly, the number of breaches that we've seen in recent years, including a very large dataset such as Equifax, suggests to me that there's a lot more work that needs to be done in general for cybersecurity and data protection.

Ms. KELLY. And also, in your written testimony, you also outlined different types of safeguards that could improve the level of protection of consumers' privacy when their data is collected and stored in AI systems. One of those safeguards is the use of a technical approach you referred to as differential privacy. Can you explain that in laymen's terms?

Mr. BUCHANAN. Sure. Simplifying a fair amount here, differential privacy is the notion that before we put data into a big database from an individual person, we add a little bit of statistical noise to that data, and that obscures what data comes from which person, and, in fact, it obscures the records of any individual person, but it preserves the validity of the data in the aggregate.

So you can imagine, if we asked every Member of Congress, have you committed a crime, most Congress people and most people don't want to answer that question. But if we said to them, flip a coin before you answer; if it's heads, answer truthfully; if it's tails, don't answer truthfully; flip another coin and use a second coin flip

to determine your made-up answer, we're adding a little bit of noise when we collect the answers at the end. And using a little bit of math at the back end, we can subtract that noise and get a very good aggregate picture without knowing which Members of Congress committed crimes.

So the broader principle certainly holds, again, with a fair more math involved, that we can get big picture views without sacrificing the privacy or criminal records of individual members of the dataset.

Ms. KELLY. I have not committed a crime, by the way.

Mr. HURD. Neither have I.

Ms. KELLY. Do you feel like if more businesses adopted this differential privacy, this type of security measure would be more effective in mitigating the risk to personal privacy?

Mr. BUCHANAN. With something like a differential privacy, the devil's in the details; it has to be implemented well. But as a general principle, yes, I think it's a very positive technical development and one that is fairly recent. So we have a lot of work to do, but it shows enormous promise, in my view.

Ms. KELLY. Thank you. And in addition to this, you also identify in your written testimony another type of security control known as on-device processing. Can you, again in laymen's terms, explain on-device processing and how it operates to protect sensitive and personal data?

Mr. BUCHANAN. Sure. This one's much more straightforward. Essentially, the notion that if we're going to have a user interact with an AI system, it is better to bring the AI system to them, rather than bring their data to some central repository. So if an AI system is going to be on your telephone—your cell phone, rather, you can interact with the system and do the processing on your own device rather than at a central server where the data is aggregated. Again, as a general principle, I think that increases privacy.

Ms. KELLY. And in your assessment, what are the reasons why more companies in general are not deploying these types of security controls?

Mr. BUCHANAN. Certainly, as a matter of practice, they require enormous technical skill to implement. Frankly, I think some companies want to have the data, want to aggregate the data and see the data, and that's part of their business model. And that's the incentive for those companies not to pursue these approaches.

Ms. KELLY. What recommendations would you have for ways in which Congress can encourage AI companies to adopt more stringent safeguards for protecting personal data from consumers?

Mr. BUCHANAN. I think Mr. Clark has made excellent points about the importance of measurement, and I think this is an area that I would like to know more about and measure better of how are American companies storing, securing, and processing the data on Americans. So that would be, Chairman Hurd mentioned measurement is a topic of interest to this committee, and I think that would be one place to start.

Ms. KELLY. And just lastly, a part of the struggles that companies have is that because they don't have enough of the expertise because it is not in the workforce?

Mr. BUCHANAN. Yes, Congresswoman, I think that's right. There's enormous demand that has not yet been met for folks with the skills required to build and secure these systems. That's true in AI, and that's true in cybersecurity generally.

Ms. KELLY. And would the rest of the witnesses agree with that also?

Mr. SHAPIRO. Yes, the last comment.

Mr. CLARK. Yes. We need 10, 100 times more people with these skills.

Ms. LYONS. I would agree with Dr. Buchanan.

Ms. KELLY. Thank you. And thank you, Mr. Chair.

Mr. HURD. Thank you, ranking member.

Mr. Shapiro, you know, I had the fortune of attending CES, the Consumer Electronics Show, this recent January. Thanks for putting on such a good show. I learned a lot about artificial intelligence and how important data is in training, the actual training the algorithm.

And one of the questions that we have come up—or we have heard on the issues we have heard is the importance of data, and we've learned about bias and preventing that. We learned about being auditable. We know we have to invest more money in AI. We also know when you train people better.

Who should be taking the lead? Like, who is the person, who should be driving kind of this conversation? Or maybe let me narrow the question. Who in government should be driving kind of the investment dollars in this? And I know you have peer research at universities. You have the national labs. You know, who should be coming up with that with our investment plan in AI?

Mr. SHAPIRO. Well, I think, first of all, we have to agree on the goals. I liked the idea of measurement as well, and I think the goals are, number one, we would like the U.S. to be the leader; two, we want to solve some fundamental human problems involving healthcare, safety, cybersecurity. Now, there's some—we can define goals with those. And third, we want to respect people's privacy.

And I think there has to be national discussion on some of these issues because let's take the issue of privacy, for example, and we've heard a lot about that today. The reality is, culturally, we're different on privacy than other parts of the world. In China, the concept of privacy is, especially in this area, is that the citizens really don't have any. They're getting social scores. Their government is monitoring what they do socially, and certainly there doesn't seem to be much legal restriction on accessing whatever people do.

Europe has taken a very different—they're really focused on privacy. They have the right to be forgotten. They have the right to erase history, something that seems an anathema to us. How could you change the facts and take them off the internet? And they've really clamped down, and they're going forward in a very arguably bold and unfortunate way on this GDPR, which is really you could argue is for privacy or you could argue is to shut Europe out in a competitive fashion.

When I look at the U.S. and our focus on innovation and our success and I compare it to Europe, I see they have maybe a handful of unicorns, you know, billion dollar valuation companies, and the

U.S. has most of them in the world, over 150. And why is that? There's many answers. It's our First Amendment, it's our diversity, it's our innovation. It's the culture we have of questioning. There's many things go to it, but part of it, we're a little more willing to take some risks in areas like exchange of information.

Europe is going forward with GDPR, and frankly, it's going to hurt American companies. I was just in France last week. It's going to hurt European companies. They're terrified of it. They're talking about trying to delay it. But it's also going to kill people, because if you can't transfer, for example, medical information from one hospital to another in the same region, that has life consequences.

So when we talk about the issue of privacy and who should lead on it, I think we should do it in a commonsense way, and we shouldn't let HIPAA, for example, be our model. The model should be what is going to be—what kind of information is warranted in this situation. We've done a lot of research and we have found, for example, that Americans are willing to give up personal information for a greater good, as they have done with health information on our Apple watches. They're willing to do it for safety. They're willing to do it for the security of their children. They're willing to do it for their own safety involving, for example, where your car is if it hits a car in front of them.

So in the privacy area, I think we have a pretty good cultural sense. I think the Federal Trade Commission has a broad mandate to do a pretty good job in that area.

And I don't want to take all the time, but those other two areas I talked about in terms of the measurements and artificial intelligence and what they should do, it goes into how you get the skilled people, what you do, how you change your educational system, how you retrain for jobs. There's a lot of things that government can do and Congress can do. And I applaud you and this committee for taking the first big step in having hearings to raise the issues, but what I would expect Congress to do in the future is rather than come up with immediate solutions, is instead to focus on what the goals are and how we could do that.

And I would look at two examples that I was both personally involved with, which was government setting big goals, but working with industry who came up with private things. Actually, I'll give three quickly.

One is, and Congressman Issa is very aware of this because he was part of it, the transition to high definition television. That was we set the goal. We wanted to have the best system in the world, private industry, no spending of government money, we did it.

Second is commercialization of the internet, doing business over it. We have done it in Virginia with bipartisan way and the goals were there and it worked.

And the third is you talked about privacy for wearable devices, healthcare devices which came up earlier. At CTA, we got every everyone in the room that made those devices and we agreed on a regimen of saying this is what we should voluntarily do. This is what we should follow. It should be transparent, clear language, opt out, and you can't sell information or use it without any permission from your customers. And the Obama administration

seemed pretty happy with that, and even they didn't act because that was industry self-regulation.

Mr. HURD. Got you. Thank you, Mr. Shapiro.

I'm going to come back for another round of questions, but now I'd like to recognize my friend from the Commonwealth of Virginia for his first round of questions.

Mr. CONNOLLY. Thank you, Mr. Chairman.

And, Gary, you were doing speed dating there. And welcome. Good to see you again.

I want to give you a little bit more opportunity maybe those three things you were just talking about if you want to elaborate a little bit more. Because this idea—to let you catch your breath for a second. This idea of the zone of privacy and some of it is cultural bound I think is absolutely true, but I can remember going to Silicon Valley about 9 years ago, meeting with Facebook people. And their view about privacy was we, Americans, need to get used to shrinking boundaries for privacy and that the younger generations were already there. Older generations needed to just learn to suck it up and accept it.

And I think watching what happened to Mr. Zuckerberg here in the last couple weeks, one needs to not be so facile. You're not being facile, but I mean, I think you're rising, though, those questions. Some of its cultural bounds, some of it the rules of engagement aren't quite there yet. We debate, do we get involved? If so, what do we do?

And so I think your thoughts are very helpful, given your experience and your position in providing some guidance. So I want to give you an opportunity to elaborate just a little bit.

Mr. SHAPIRO. Well, thank you very much, Congressman. I appreciate it. I guess my view is that as a Nation, we're not China where we totally don't devalue privacy, and we're not Europe where we use privacy as a competitive tool against other countries, frankly, but it also tamps down innovation in our own country.

Our competitive strength is innovation. That's what we're really good at. It's the nature of who we are. So the question is, how do we foster innovation in the future in AI and other areas and also maintain our—correct our citizens' view that they are entitled to certain things?

Now, to a certain extent, it's educating. Everyone has an obligation. The obligation of business is to tell our customers what it is we're doing with their data in a clear and transparent way, and frankly, we haven't done a great job at it. I mean, if I had my way, I wouldn't want to have to click on those "I agree" just to get to the website I want to. I'd like to click on platinum, gold, or silver standard. If there's some standardization, it would probably help, and government can play a role in that.

But we also want to make sure that we can innovate. And consumers should understand that they're giving away something in return for free services. You give a tailor your information on your body size to get clothes that fit. You give your doctor information about your health, and you're always giving away something. And, you know, the truth is if you're going to get a free service, like Facebook or Google, and you want to keep it free, they are using that data to get people to know you.

But it's like I shop at Kroger's in Michigan actually, because that's where I commute to, and Kroger's knows a lot about me. They know everything I buy, and they give me coupons all the time. And I value those coupons. But they know what I buy, and I am willing to do that. It's the same thing with other frequent user programs. We're doing that all the time. We're giving up information about ourselves. We get discounts, we get deals, and we get better service.

If we do it with our eyes open and we're educated about it, that's fine. Now, the role of citizens is to understand——

Mr. CONNOLLY. By the way, we know you shopped at Kroger's last Thursday, and that fondness you've got for frozen rolls has, frankly, surprised us.

Mr. SHAPIRO. So in terms of the role of government, I think the role of government is to start out by having hearings like this one, define the goals and the measurements culturally for the future. And the role, frankly, of the administration, in my view, is to set the big goals and to make sure that we buy into them on a bipartisan way. And I love the idea of some big goals, as Mr. Clark suggested, because we need big goals in this area.

You know, for example, having self-driving cars by 2025 or nothing—dropping the death rate from automobiles down by half by a certain date would be a very admirable goal that everyone in this country can rally around.

Mr. CONNOLLY. Thank you so much, Gary.

Mr. Clark, in the time I've got left, you said in the report on OpenAI, that artificial intelligence continues to grow. Cyber attacks will utilize AI and will be, and you said, quote, "more effective, finely targeted, difficult to attribute, and likely to exploit vulnerabilities in AI systems."

I want to give you an opportunity to expand a little bit on that. So how worried should we be?

Mr. CLARK. So you can think of AI as something that we're going add to pretty much every aspect of technology, and it's going to make it more powerful and more capable. So this means that our defenses are also going to get substantially better. And as Dr. Buchanan said earlier, you weren't in the room, it's not clear yet whether this favors the defender or the attacker. And this is why I think that hosting competitions, having government measure these capabilities as they develop, will give us a kind of early warning system.

You know, if there's something really bad that's about to happen as a consequence of an AI capability, I'd like to know about it, and I'd like an organization or an agency to be telling us about that. So you can think about that and take that and view it as an opportunity, because it's an opportunity for us to learn in an unprecedented way about the future before it happens and make the appropriate regulations before harm occurs.

Mr. CONNOLLY. If the chair will allow, I don't know if Dr. Buchanan or Ms. Lyons want to add to that, and my time is up.

Ms. LYONS. I have nothing more to add. Thank you.

Mr. BUCHANAN. I think we probably can return to the subject later, but I would suggest we have seen some indications already

of increased autonomy in cyber attack capabilities. There's no doubt in my mind we will see more of that in the future.

Mr. HURD. The distinguished gentleman from California is now recognized for his round of questions.

Mr. ISSA. You know, this is what happens when you announce your retirement, you become distinguished.

You know, I know in these hearings that there's sort of an exhaustive repeat of a lot of things, but let me skip to something I think hasn't happened, and I'll share it with each of you, but I'll start with Mr. Clark.

The weaponization of artificial intelligence, there's been some discussion about how far it's gone, but it's inevitable. The tools of artificial intelligence disproportionately favor U.S. companies.

Now, when that happened in satellites, nuclear capability, and a myriad of data processing, we put stringent export control procedures on those things which may have a dual use. We've done no such thing in artificial intelligence. Would you say today that that is an area in which the Commerce Department's export assistant secretary doesn't have specific authority but needs it?

Mr. CLARK. Thank you. I think this is a question of existential importance to, basically, the world. The issue with AI is that it runs on consumer hardware, it's embodied in software, it's based on math that you can learn in high school. You can't really regulate a lot of aspects of fundamental AI development because it comes from technology which 17-year-olds are taught in every country of the world, and every country is developing this. So while the U.S. economy favors the development of AI here and we have certain advantages, other countries are working on this.

So I think for export controls, arms controls, do not really apply here. We're in a new kind of regime, because you can't control a specific thing with this AI technology. Instead, you need to develop norms around what is acceptable. You need to develop shared norms around what we think of an AI as safety, which is about being able to offer guarantees about how the systems work and how they behave, and we need to track those capabilities.

So I think that your question's a really important one, and I think it touches an area where much more work needs to be done because we don't have the right tool today to let us approach the problem.

Mr. ISSA. And let me follow up quickly. When we look at artificial intelligence, we look at those producing advanced algorithms. And I went to a different high school apparently than you did. Mine wasn't Caltech. So let's assume for a moment that it's slightly above high school level. The creators of those, and let's assume for a moment, hypothetically, they're all in the first world, and the first world defined as those who want to play nice in the sandbox: you, us, Europe, and a number of other countries.

Do you believe, if that's the case, the government has a role, though, in ensuring that when you make the tool that is that powerful, the tools that, if you will allow it to be safely controlled, are also part of the algorithm? In other words, the person who can make a powerful tool for artificial intelligence also can, in fact, design the safety mechanism to ensure that it wouldn't—couldn't be

used clandestinely. Do you think that's a social responsibility of, let's say, the Facebooks and the Googles?

Mr. CLARK. I think we have a social responsibility to ensure that our tools are safe and that we're developing technologies relating to safety and reliability in lockstep with capabilities. You know, that's something that the organization I work for, OpenAI, does. We have a dedicated safety research team, as does Google, Google's DeepMind, they do as well. So you need to develop that.

But I think to your question is how do you, if you have those tools, make sure everyone uses it? I think there you're going to deal with kind of two stages.

Mr. ISSA. As we've discovered today that we've sent our CIA director to meet with Kim Jong-un because he can't be trusted with the tools he's created that got to him. I might have a different view on the export controls.

But, Mr. Shapiro, since you've given me every possible look on your very creative face as these answers came through, let me allow you to answer that, but I want to shift to one other question. You mentioned a little bit HIPAA. Now, the history of HIPAA is precomputer data. It is, in fact, a time in which, basically, pieces of paper were locked up at night and not left out on desks so that one patient didn't see another patient's records and that you didn't arbitrarily just answer anyone over the phone.

The reality, though, today is that the tools that your industry, the industry you so well represent, you have tens of thousands of tools that are available that can gather information, and often, they're limited by these requirements from really interacting with the healthcare community in an efficient way. Do we need to set up those tools to allow healthcare to prosper in an interactive cloud-based computer generation?

And I'll just mention, for example, the problem of interoperability between Department of Defense, the Veterans Administration, and private doctors, that has been one of the—and it's confounded our veterans, often leading to death to overdose for a lack of that capability. Do you have the tools, and what do we need to give you to use those tools?

Mr. SHAPIRO. Well, it's probably fair to say that the promise of ObamaCare, which was very positive of allowing easy transfer of electronic medical records, has not been realized. I think even the American Medical Association, which urged that it be passed, has now acknowledged that, and it's been a great frustration to doctors, as I think you know.

In terms of the tools that we have today to allow easy transfer, you know, the administration hasn't endorsed this blue button initiative which allows medical records, especially in emergency cases, to be transferred easily. I think we have a long way to go as a country to make it easy to transfer your own health information. The old ways they did it in the communist countries is you walk around with your own records. Your paper records were actually a simpler transaction than what we have today where everyone goes in and has to start from zero.

Mr. ISSA. Well, you know, the chairman is too young to know, but I walked around in the Army with that brown folder with all my medical records, and it was very efficient.

Mr. HURD. What is a folder?

Mr. ISSA. What's a folder?

Mr. SHAPIRO. But the opportunity we see and we're concerned as an organization about the growing deficit and the impact that will have existentially on our country, frankly, and we see the opportunity there in providing healthcare and using technology in a multitude of ways to lower costs, to be more efficient, to cut down on doctor visits, and to just allow easy transfer of information.

In terms of what the government can do, we're actively advocating for a number of things. We're working with the FDA. We're moving things along. And we found with this administration and prior administration a great willingness to adopt the technology system; it is just a matter of how fast.

Mr. ISSA. Thank you. Mr. Chairman, are we going to have another round?

Mr. HURD. Yes.

Mr. ISSA. Okay. I'll wait. Thank you.

Mr. HURD. Ranking member, round two.

Ms. KELLY. Thank you.

Given Facebook CEO Mark Zuckerberg's comments last week to Congress, how would you evaluate AI's ability to thwart crime online, from stopping rogue pharmacies, sex trafficking, IP theft, identity theft to cyber attacks? And whoever wants to answer that question I'm listening.

Mr. BUCHANAN. I think, speaking generally here, there's enormous promise from AI in a lot of those areas, but as I said in my opening remarks, we should recognize that technology will not replace policy. And I think it's almost become a cliché in certain circles to suggest that, well, we had this very thorny, complex interdisciplinary problem so let's just throw machine learning at it and the problem will go away. And I think that's a little bit too reductive as a matter of policymaking.

Ms. KELLY. Anybody else?

Ms. LYONS. I would echo Dr. Buchanan's remarks, insofar as I think part of the solution really needs to be in bringing multiple solutions together. So I think policy is certainly part of the answer. I think technology and further research in certain areas related to security, as you mention, in the specific case is the answer. And I think also, you know, that is sort of the project of the organization that I represent, insofar as the interest of bringing different sectors together to discuss the means by which we do these things in the right ways.

Ms. KELLY. Thank you.

Dr. Buchanan, what types of jobs do you see that will be threatened in the short term by AI automation, and what about in the long term as well?

Mr. BUCHANAN. Certainly, in the near term, I think the jobs that are most at risk are jobs that involve repetitive tasks, and certainly this has always been the case with automation. But I think, as you can imagine, as artificial intelligence systems become more capable, what they can do, what they consider repetitive certainly would increase. And I think jobs that involve, again, repetition of particular tasks that are somewhat by rote, even if they're jobs

that still involve intellectual fire power are on balance more likely to be under a threat first.

Ms. KELLY. And in the long term what do you see?

Mr. BUCHANAN. As we move towards an era of things like self-driving cars, one could imagine that services like Uber and Lyft might not see a need for drivers and taxis, might not see a need for—taxi companies, rather, might not see a need for drivers. There's some suggestion that if we had such a world, we would need fewer cars in general. Certainly, Members of Congress are acutely aware of how important the auto industry is in the United States.

So when you look at a longer term horizon, I think there's more uncertainty, but there's also a lot more potential for disruption, particularly with knock-on effects of if the auto industry is smaller, for example, what would the knock-on effects be on suppliers to companies even beyond just the narrow car companies themselves.

Ms. KELLY. And to whoever wants to answer, what type of investments do you feel that we should be making now for people that are going to probably lose their job? What do you—how do you see them transitioning to these type of jobs?

Of course.

Mr. SHAPIRO. So I would—the prior question about the jobs. The great news is the really unpleasant, unsafe jobs, most of them will go away. So, for example, I was using a robotics company, and they have something specialized which is very good at picking up and identifying and moving things around using AI and robotics. The jobs—one of the potential buyers was a major pizza company that delivers to homes. The way they do it is they make dough, and the dough today is made by people in a very cold, sterile environment. They wear all this equipment to be warm and also to be sterile. And they can only work—it's very ineffective. No one wants to do the job at all. And this solves that problem.

There's also, you know, thousands of other conditions where jobs are really difficult. It could be picking agriculturally where now there's—increasingly there's devices which do that, and they do have to be fairly smart to identify the good versus the bad and what to pick and what not to pick.

In terms of what investment has to make in terms of retraining. I think we have to look at the best practices in retraining and figure out what you could do. I mean, we do have millions of unemployed people today, but we have more millions of jobs that are open and not filled. Some of it is geographic, and we should be honest about. And maybe we need to offer, you know, incentives for people to move elsewhere.

But some of it is skills. And the question is what has worked before, what skills that you can train someone for, whether it's a customer service center or whether it's something basic programming or helping out in other ways.

I think we have to look at individual situations, ferret out what's out there that's already worked and try some new things, because a lot of what has worked in the past will not work in the future. And the longer term investment is obviously with our kids. We just have to start training them differently.

And we also have to bring back respectability to work which is technical work, as Germany has done, and focus on apprentice programs and things like that, and not just assume that a 4-year degree is for every American, because it's not a good investment of society. And there's a lot of unemployed people who went to college who don't have marketable skills.

Ms. KELLY. Mr. Clark.

Mr. CLARK. So I think that this touches on a pretty important question which is, where the job gets created, because new jobs will be created, will be uneven. And where the jobs get taken away will also be uneven.

So I want to refer to a couple of things I think I've already mentioned. One is measurement. It's very difficult for me to tell you today what happens if I drop a new industrial robot into a manufacturing region. I don't have a good economic model to tell you how many jobs get lost, though I have an intuition some do. That's because we haven't done the work to make those predictions.

And if you can't make those predictions, then you can't invest appropriately in retraining in areas where it's actually going to make a huge difference. So I want to again stress the importance of that measurements and forecasting role so the government can be effective here.

Ms. KELLY. Thank you very much. I yield back.

Mr. HURD. Mr. Clark, you talked about a competition, you know, akin to robotics I, akin to self-driving car. What is the competition for AI? What is the question that we send out and say, hey, do this thing, show up on Wednesday, August 19, and bring your best neural network and machine learning algorithm?

Mr. CLARK. So I have a suggestion. The suggestion is a multitude of competitions. This being the Oversight Committee, I'd like a competition on removing waste in government, you know, bureaucracy, which is something that I'm sure that everyone here has a feeling about. But I think that that actually applies to every committee and every agency.

You know, the veterans agency can do work on healthcare. They can do a healthcare moonshot within that system that they have to provide healthcare to a large number of our veterans. The EPA can do important competitions on predicting things like the environmental declines in certain areas affected adversely by extreme weather.

Every single agency has data. It has intuitions of problems it's going to encounter and has competitions that it can create to spur innovation. So it's not one single moonshot; it's a whole bunch of them. And I think every part of government can contribute here, because the great thing about government is you have lots of experience with things that typical people don't. You have lots of awareness of things that are threats or opportunities that may not be obvious. And if you can galvanize kind of investment and galvanize competition there, it can be kind of fun, and we can do good work.

Mr. HURD. So along those lines, how would you declare a winner?

Mr. CLARK. In which aspect?

Mr. HURD. Let's say we were able to get—well, we can take—let's take HHS. Let's take Medicare. Medicare overpayments. Perfect example. And let's say we were able to get that data protected in a

way that the contestants would be able to have access to it. And you got 50 teams that come in and solve this problem. How would you grade them?

Mr. CLARK. So NAI, we have a term called the objective function. What it really means is just the goal. And whatever you optimize the goal of a thing for is what you'll get out. So doing a goal selection is important because you don't want to pick the wrong goal, because then you'll kind of mindlessly work towards that.

But a suggestion I'd have for you is the time it takes a person to flow through that system. And you can evaluate how the application of new technologies can reduce the time it takes for that person to be processed by the system, and then you can implement systems which dramatically reduce that amount of time. I think that's the sort of thing which people naturally approve of.

And just thinking through it on the spot here, I can't think of anything too bad that would happen if you did that, but I would encourage you to measure and analyze it before you set the goal.

Mr. HURD. Ms. Lyons, what's the next milestone when it comes to artificial intelligence?

Ms. LYONS. From a technical perspective or otherwise?

Mr. HURD. From a technical perspective.

Ms. LYONS. Well, I think what a lot of the AI research community is looking towards is AI platforms that can be applied to more generalized tasks than just the very narrow AI that we see applied in most of the circumstances that we've described today.

So I would say that's—that is the next sort of moonshot milestone for the technical community.

Mr. HURD. Is that a decade? Is that 9 months? Is it 20 years?

Ms. LYONS. You know, I have my own personal perspectives on this. The Partnership hasn't really formulated one yet. But I think we have a lot of organizations involved in ours which have disagreeing viewpoints on this. And I'm sure, actually, if this committee was quizzed, we might all have different answers as well.

But I think we are—we're years and years away from that. And it's useful to be thinking about it right now, but I do think we're probably decades.

Mr. HURD. And what are the elements that are preventing us from getting there?

Ms. LYONS. I actually don't think I'm the best person equipped on this panel to give you an answer to that. I'm pretty far away from the technical research, from where I'm sitting right now. But it is a—it is a—they are technical impediments that are stopping us from achieving that at this moment.

Mr. HURD. Good copy.

Dr. Buchanan, how do we detect bias?

You know, I think one of the things that we have heard through these hearings is bias. And we know how you create bias, right. Giving—not giving a full dataset, right? So you can—can the algorithm itself be biased? Is the only way to introduce bias is by the dataset? And then how are we detecting whether or not the decisions that are being made by the algorithm show bias?

Mr. BUCHANAN. I'm not convinced that we're looking as much as we should. So when you say how are we detecting, I think in many cases we are not detecting bias systems.

But speaking generally of how do you—

Mr. HURD. Philosophically, how do we solve that problem?

Mr. BUCHANAN. Right. I think it's worth—again, as I said before, technology cannot replace policy, and we should first develop an understanding of what we mean by a bias. Is a system biased, whether it's automated or not, if it disproportionately affects a particular racial group or gender or socioeconomic status? And I think that most people would answer yes. And you would want to look at what the outcomes of that system were and how it treated individuals from certain groups.

And there's a number of different values you can instantiate in the system that try to mitigate that bias. But bias is a concept that we intuitively all feel, but it's often quite difficult to define. And I think a lot of the work in detecting bias is first work in defining bias.

Mr. HURD. Mr. Clark.

Mr. CLARK. So I have two suggestions. I think both are pretty simple and doable. One is whenever you deploy AI, you deploy it into a cohort of people. Let's say I'm deploying a public service speech recognition system to do a kind of better version of a 311 service for a city. Well, I know I have demographic data for that city and I know that people that speak, perhaps not with my accent, but a more traditional American one are going to be well represented.

Mr. HURD. Do you have an accent?

Mr. CLARK. It's sometimes called Australian, but it's actually English.

So I think that, when you look at your city, you're going to see people who are represented in that city but are not the majority. So you test your system against the least represented people and see how it rates. That will almost invariably surface areas where it can be improved.

And the second aspect is you need more people in the room. This requires like a concerted effort on STEM education and on fixing the diversity in STEM, because if you're not in the room, you probably just won't realize that a certain bias is going to be obvious, and we do need to fix that as well.

Mr. HURD. So we've all—and—oh, Ms. Lyons, go right ahead.

Ms. LYONS. I might just chime in by summarizing, because I don't think I heard this clarified in the way that I might describe it, which is in the different ways in which bias is represented in systems. And I think that is through data inputs, which we've talked a little bit about. It's also in the algorithms themselves. And I think that gets to some of the points that Mr. Clark has made around who is building the systems and how representatives those developer communities are.

And then I think further than that, it's also in the outcomes represented by those various inputs and the ways in which there might be adverse or outsized impacts on particularly at-risk communities who are not involved in technology development. So I just wanted to add that.

Mr. HURD. That's incredibly helpful.

We've all been talking about what should the standards be or what is the—what are the—what are the equivalent of the three rules from I, Robot, right?

And in the one that it seems that there's most agreement, and correct me if I'm wrong on this, is making sure the decisions of the algorithm are audible, that you have—that you understand how that decision was made by that algorithm. There's been so many examples of an AI system producing something, and the people that design the algorithm have no clue how the algorithm produced that.

Is that the first rule of artificial intelligence? What are some potential contenders for the rules of ethical AI?

And, Dr. Buchanan, maybe start with you, go down the line, if anybody has opinions.

Mr. BUCHANAN. I suggest that the first rule might generate more discussion than you'd expect on this panel.

In general, there is oftentimes a tradeoff because of the technology involved in AI systems between what we call the explainability or interpretability of an algorithm's decision and how effective the algorithm is or how scaleable the algorithm is.

So while I certainly think it's an excellent aspiration to have an explanation in all cases, and while I probably believe that more than many others, I could imagine cases in which we worry less about how the explanation—or how the algorithm makes its decision and more about the decision.

For example, in medicine, we might not care how it determines a cancer diagnosis as long as it does so very well. In general, however, I suggest explanations are vitally important, particularly when it comes to matters of bias and particularly given the technology involved, they're often hard to get.

Mr. HURD. Anybody else have an opinion?

Ms. Lyons.

Ms. LYONS. I think the question that you've raised is actually fairly central to ongoing dialogues happening in the field right now. And the easy answer is that there is no easy answer, I think. And I think Dr. Buchanan has demonstrated that with his remarks as well.

But generally speaking, I do think that it's—it has been a particular focus, especially of—especially in the last several years, of a certain subset of the AI machine learning technical community to consider questions associated with issues regarding fairness, accountability, transparency, explainability. And those are issues associated with auditability, as you describe. And a keen interest, I think, in making sure that those conversations are multidisciplinary in nature as well, and including people from fields, not necessarily traditionally associated with computer science and AI and machine learning communities, but also inclusive of the ethics community, law and policy community, and the sociology community more generally. So—

Mr. HURD. So are there other things—I recognize that this is a loaded question and there's not an agreement on this. But what are some of the other contenders that people say, hey, we should be doing this, even if we don't—even if we recognize there's not agreement, when it comes to what are the rules of ethical AI?

Ms. LYONS. Well, the Partnership, for its part, has a set of tenets which are essentially eight rules that govern the behavior at a very broad level of the organizations associated with us. And they're posted on our website. We included them in our written remarks—or written testimony as well.

But generally speaking, at a high level, we have, as a community, decided on certain sort of codes of conduct to which we ascribe as organizations involved in this endeavor. And I think a central project of this organization moving forward from this point is in figuring out how to actually operationalize those tenets in such a way that they can be practiced on the ground by developers and by other associated organizations in the AI technical community.

Mr. HURD. Mr. Clark.

Mr. CLARK. So I want to put a slightly different spin on this, and that's about making sure that the decisions an AI makes are sensible. And what I mean by sensible is, you know, why do we send people to school? Why do we do professional accreditation? Well, it's because we train people in a specific skill, and then they're going to be put into a situation that they may not have encountered before, but we trust with the training and education they had in school means that they'll take the correct action. And this is a very common thing, especially in areas like disaster response where we train people to be able to improvise. And these people may not be fully auditable, like you ask him why did you do that in that situation? And they'll say, well, it seemed like the right thing to do. That's not a super auditable response, but it's because we're comforted in the training they've had.

And so I think some of it is about how do we make sure that the AI systems we're creating are trained or taught by appropriate people. And that way we can have them act autonomously in ways that may not be traditionally interpretable, but we'll at least say, well, sure, that's sensible, and I understand why we trained them in that way.

Mr. HURD. Mr. Shapiro, you have 45 seconds, if you'd like to respond.

Mr. SHAPIRO. You've raised so many issues that I'll pass on this one.

Mr. HURD. Mr. Issa, you're now recognized for round two.

Mr. ISSA. Thank you.

While you were doing that, I was listening to the deactivation of HAL 9000.

[Audio recording played.]

Mr. ISSA. Well, it's not 2001 anymore, but it is.

You know, this dialogue on AI, this portion of it I think is, particularly for people who saw that movie, knew is important because HAL was able to correspond, able to have a dialogue, but it didn't have to answer honestly, and it didn't have to be, if you will, proofed. In other words, nobody put in the ability in the algorithm for it to be queried and to answer.

So I think for all of us that are having this dialogue and for those of you working in it, the question is will we make sure that we have open algorithms, ones that can be queried and, as a result, can be diagnosed. And if we don't have them, then what you have

to rely on, as was being said, is outcome. Outcome is not acceptable.

Outcome is why the IRS shut down yesterday and wasn't taking your tax returns on the tax day, and all they knew was they had to fix it, but they didn't know why it happened, or at least it happened in a portion of their system.

So that's something that I can't do. We can't mandate. But there are some things that, hopefully, we can work on jointly.

And, Mr. Shapiro, you know, nearly 100 years ago, the Radio Manufacturers Association formed. And one of things it began to do was standard setting. Earlier today, you talked about we should have a standard, if you will, for what am I disclosing? Platinum, gold, silver. You had a way of saying it.

My question to you is, where are the responsible parties, as the radio manufacturers, now CTA, 100 years ago, who began saying, if we're going to have the flourishing of technology, we're going to have to have standards? Privacy standards are complex, but how do you make them simple? Well, you make them simple by building standards that are predictable that people can share a decision process with their friends. Yes, I always go for silver if it's my medical and gold if it's my financial.

You alluded to it. How do we get there knowing it's not going to be mandated from this side of the dais? Or at least we certainly couldn't come up with the examples.

Mr. SHAPIRO. Well, I mean, that's not the only choice. I mean, there's the executive branch. The FTC is comfortable in that area. And sometimes—

Mr. ISSA. But wait a second. I know the FTC. They're very comfortable, after something goes bad, telling us it went wrong.

How often do they actually predictively able to say what the, quote, industry standard is before something? They certainly haven't done it in data intrusions.

Mr. SHAPIRO. Well, they do have a history of providing guidelines and standards. And I'm not advocating that. I'm not—what I'm saying is, on the issue of privacy and click-on, there are so many different systems out there that I am not personally convinced that the industry could come forward together without some concern the government would instead. I think it's always preferable for government and industry to work together, but sometimes the concern that government will act does drive industry to act. That's just the reality.

In this area, it's—that cat's out of the bag a long time ago, and we're all clicking on stuff we don't understand. And that may have been one of the issues, even in the Facebook disclosures and things like that, which I think cause some concern, is that we're agreeing on things that we don't understand. I mean, I used to read that stuff. I've stopped a long time ago. It's just—you can't read it or understand it.

Mr. ISSA. But, Gary, back to what we were talking about in the last round. When we look at the healthcare and at personal information related to your healthcare, your drugs, your body weight, whatever it is, those are not such a large and complex hypothetical. Those are fairly definable.

If we want to have the benefits of group data, such as a Fit Bit gives and other data, and yet protect individual privacy, isn't this a standard that we should be able to demand be produced and then codified hopefully with some part? I mean, the FTC is very good if you produce a standard of saying that's now the industry standard. They're less good at defining it and then—proactively.

Mr. SHAPIRO. Well, thank you for raising that specific case. We have done that as an industry. We've come up with standards. They are voluntary, and we haven't heard about any data breach, that I'm aware of, in the personal wearable area, because I think that was a model that came together.

The automobile industry is doing something similar, and other industries are doing it. It's not too late. I was just talking about the online click agreements.

Mr. ISSA. Sure.

Mr. SHAPIRO. There's opportunity in other areas. And I think to move forward and to move forward quick, it's an opportunity. The advantage for the companies is they're kind of safe in the herd if they follow the herd.

Mr. ISSA. Wait. And I'm going to cut you off but not finish this.

One door down there in Judiciary, we control the question of limiting liability for certain behavior or not limiting it. Where, in your opinion—and I can take the rest of you if the chairman will allow—where is it we need to act to show that if you live by those best practices, knowing that, just like that thing I played, it will not be 100 percent. But if you live by those practices, your liability is in some way limited. In other words, nonpunitive if you're doing the right things. Because right now, Congress has not acted fully to protect those who would like to enter this industry.

Mr. SHAPIRO. Well, you've asked the question and answered it at the same time. Obviously, being in business yourself, you understand that risk and uncertainty are factors. We're seeing that in the trade problems we face today, the potential tariffs that—

Mr. ISSA. I did have that told to me just last night by somebody who knows about the question of not setting prices for their customers for Christmas because they don't yet know what the tariff will be.

Mr. SHAPIRO. So uncertainty in the business environment. We're seeing it increasingly reflect in the stock market. But in terms of potential liability, our companies welcome certainty.

And one thing, for example, Congress did when credit cards were introduced, they said your liability as an individual is limited to \$50, and it all of a sudden allowed people to get over that uncertainty of going from cash to credit cards. And it helped grow our economy enormously and take a lot of friction out.

We're facing some of the same things now as we go forward in so many different areas because of AI. And we do have an opportunity for Congress to address and say, if you follow these practices, you have a safe harbor here. But that's a very difficult thing to do, and especially when it gets to the base of our privacy and leaks and things like that.

But everyone's looking for best practices in the issues we're discussing earlier having to do with cyber and how do you protect. I mean, this is—game will never end. You build a better mousetrap,

you get smarter mice. So we're going to keep raising that bar, and that's the challenge that Congress will face.

But some safe harbors would be certainly welcome in this area as grows rapidly. And I think there's a role to play. And I think this is a great amazing set of first three hearings to start on what will be a process with government and industry and consumers.

Mr. ISSA. I hear that from all of you. I saw a lot of heads nodding, that safe harbors should exist if we're going to promote the advancement of and use of data in our artificial intelligence.

Any noes?

Ms. LYONS. Well, I—

Mr. ISSA. There's always a caveat, but any noes?

Ms. LYONS. I actually don't—I don't really have any comments about safe harbors specifically. But I think, in general, the issue of generating best practices is one which is really important to be considered in this field. And that, again, was sort of the reason why the Partnership on Artificial Intelligence was created, because there is a sort of understanding, I think, that's been come to in a collective sense about the necessity of determining what these guardrails should be, to a certain extent. And I think that project can't really be undertaken without the policy community as well as other stakeholders who just necessarily need to be involved.

Mr. ISSA. Thank you, Mr. Chairman.

Mr. BUCHANAN. I would also put myself down as embracing a caveat here, Congressman. I think one of the dangers is that we agree on a set of best practices that are not, in fact, anywhere near best practices and we think our work is done.

So while I support safe harbors if they align to practices that do protect privacy and advance security, I would suggest we are long way from those practices in place today. So we should not lock in the status quo and think our work is done.

Mr. ISSA. Thank you.

You know, I've owned Model Ts. I've owned cars from the fifties, sixties, seventies, eighties and so on. I don't think we lock in best practices. We only lock them in for a product at the time that the product is new and innovative and we have an expectation for the manufacturer that that product will become obsolete. Nobody assumes that at a Model T is the safest vehicle, or even a '66 Mustang.

But we do we do make expectations at the time of manufacturing. You know, there was a time when, years ago, when a man lost a limb on a lathe, and he sued the company, even though the lathe had been made in 1932, and it was already, you know, 50 years later. And we had to create a law that prohibited you from going back and using today's standards against the manufacturer. You could use it against the company if they hadn't updated, but you couldn't use it against the manufacturer.

That's an example of safe harbor where, if you make to the standards of the day, you are not held for the standards that change on a product that is a fire-and-forget. You don't own it or control it. And so that's what I was referring to, your expectation that, yes, there has to be continuous innovation and that people have to stay up with the standards. Of course, we're not expecting that. But then the question is, will we see it from your side, or

would we try to have the same people, you know, who have the system that shut down on the last tax filing day be the ones determining best practices.

Thank you, Mr. Chairman.

Mr. HURD. Would the gentleman engage in a colloquy?

Mr. ISSA. Of course.

Mr. HURD. What's a Model T?

No, I'm joking.

Mr. ISSA. Well, you know, I just want you to know that when the big bang comes, the Model T is one of the vehicles that will still crank up and run.

Mr. HURD. Well, I'm coming to your house, Congressman Issa.

I have two final questions. The last question is actually a simple question. But the first question is—I recognize we can have a whole hearing on the topic. And I lump it generally in pre-crime, right? You have jails that are making decisions on whether someone should be released based on decisions based on algorithms. We have people making a decision about whether they believe someone's going to potentially commit a crime in the future. And I would lump this in pre-crime. And the question is, should that be allowed?

Gary.

Mr. SHAPIRO. I'll foolishly take a shot at that. It depends on the risk involved. For example, in an airplane security situation, I think it makes sense to use biometrics and predictive technology and gait analysis and voice analysis and all the other tools that are increasingly available to predict whether someone's a risk on a flight. Israel does it increasingly, and it's—it makes sense.

In a penal release system, I think we have more time and we are more sensitive to the fact that there are clearly racial differences in how we've approached things since day one. It may not make that much sense, so I'd say it's situational.

Mr. HURD. Mr. Clark.

Mr. CLARK. We have a robot at OpenAI, and we trained it to try and reach towards this water bottle. And so we obviously expected that the robot would eventually grab the water bottle and pick it up. But what we discovered the robot had learned to do was to take the table the water bottle was on and just bring it towards itself fulfilling the objective but not really in the way we wanted it to.

So I think I'd agree with Gary that maybe there are some specific areas where we're comfortable with certain levels of classification because the risk of getting it wrong, like a plane is so high. But I think we should be incredibly cautious, because this is a road where, once you go down it, you're dealing with people's lives. And you can't, in the case of pre-crime, really inspect wherever it's pulling that table towards it. It may be making completely bananas decisions and you're not going to have an easy way to find out, and you've dealt with someone's life in the process. So I'd urge caution here.

Ms. LYONS. I'll say this with the caveat that I provided previously on other answers, which is that the Partnership hasn't yet had a chance to formulate a position on this formally. But I think that this question speaks to a lot of the challenges associated with bias in the field right now, which we discussed a little bit earlier.

And I think also the challenges of what happens as a result of the decontextualization of technology and the application of it in areas where it may or may not be appropriate to have it be applied.

So I think it's really important to consider the impacted communities, especially in the case of criminal justice applications. And I think that needs to be a required aspect of conversation about these issues.

Mr. HURD. Dr. Buchanan.

Mr. BUCHANAN. I'd echo Ms. Lyons' points and Mr. Clark's points. I would make three other points here.

The first is that, not only is there a risk of bias, but there's a risk—sometimes machine learning is said to be money laundering for bias in that it takes a system that is something that's dirty and outputs it in this veneer of impartiality that comes from the computer. And we don't interrogate that system as much as we should. It's a major risk, I think, in this area but in many areas.

Secondly, I think you posed the question somewhat of a hypothetical. Mr. Clark is a measurer here, but I would encourage you and Mr. Clark to investigate how much the systems are already in place. I think ProPublica did an excellent bit of reporting on bias in sentencing in the criminal justice system already in place today. And that would certainly deserve more attention, in my view.

And the third is that we should make sure that the inputs to the system are transparent and the system itself is transparent. And one of my concerns, speaking generally here, is that the systems used for sentencing now and in the future often are held in a proprietary fashion. So it's very hard to interrogate them and understand how they how work. And, of course, hard in general to understand the outputs of such a system. And I think while that causes me concern in general, it should cause extreme concern in this case if we're sentencing the people on the basis of proprietary closed systems that we do not fully understand in public view.

Mr. HURD. Thank you, Dr. Buchanan.

And my last question is for the entire panel, and maybe, Dr. Buchanan, we start with you, and we'll work our way down. And it's real simple. Take 30 seconds to answer it.

What would you all like to see from this committee and Congress when it comes to artificial intelligence in the future?

Mr. BUCHANAN. Mr. Chairman, I think you've done a great job by holding this series of hearings. And I was encouraged by your suggestion that you'll produce a report on this.

I think that the more you can do to force conversations like this out in the open and elevate them as a matter of policy discourse is important. I would suggest, as an academic, I view my job to think about topics that are important but are not urgent, that are coming but are not here in the next month or two. I would suggest that many committees in Congress should take that as a mandate as well, and I would encourage you to adopt that mindset as you approach AI.

There are a lot of very important subjects in this field that will never reach the urgency of the next week or the next month, but will very quickly arrive and are still fundamentally important to virtually all of our society.

Mr. HURD. Ms. Lyons.

Ms. LYONS. At the risk of redundancy, I also want to say thank you for the engagement, Chairman. I think that having more of these types of conversations and more knowledge transfer between those working on technology and those governing it in fora like this is deeply important.

And I think—again, I'd like to offer myself and the rest of the organizations in the Partnership as a resource to whatever extent is possible in that project of education and further understanding. And I think that it's deeply important for our policymakers as well to consider the unique impact and role that they might have in technology governance, especially within the context of a multi-stakeholder setting, which is especially characteristic, I think, of the AI field right now.

Thank you.

Mr. HURD. Well, before we get to you, Mr. Clark and Mr. Shapiro, you all aren't allowed to thank us, because I want to thank you all. We have to prevent, as we've learned in the last couple of weeks, and many of our colleagues in both chambers are unfamiliar with basic things like social media, and so we have to elevate the common body of understanding on some of these topics. And so you all's participation today, you all's written statements, you all's oral arguments help inform many of us on a topic that, you know, if—when we—when I went around the streets here in the Capitol and asked everybody what is AI, most people, if they were older than me, described how—that's why I was laughing when Issa brought that in. And people that were younger than me referred to Ava, right, from *Ex Machina*. And so you all are helping to educate us.

So, Mr. Clark, Mr. Shapiro, what should this committee and Congress be doing on AI?

Mr. CLARK. Until the first time I tried to build a table, I was a measure once, cut twice, cut type of person. And then after I built that really terrible broken table, I became a measure twice, cut once person.

The reason why I say that is that I think that if Congress and the agencies start to participate in more discussions like this, and we actually come to specific things that we need to measure that we want to build around, like competitions, it will further understanding in sort of both groups. Like, there's lots that the AI community can learn from these discussions. And I think the inverse is true as well. So I'd welcome that, and I think that's probably the best next step we can take.

Mr. HURD. Mr. Shapiro, last word.

Mr. SHAPIRO. I'm happy to embrace my colleagues' offers and views and appreciation. I have three quick suggestions.

One, I think you should continue this, but go to field hearings, to great places where there is technology, like Massachusetts or Las Vegas in January, CES.

Second, I think government plays a major role, because government's a big buyer. In terms of procurement, I think you should focus on where AI could be used in procurement and set the goals and the results rather than focus on the very technical aspects of it.

Third, in terms of—I think also that—while Congress may not easily get legislation, it could have a sense of Congress. It could

add a sense of Congress that it's an important national goal that we cut automobile deaths or we do certain things by a certain date. And setting a national goal with or without the administration could be very valuable in terms of gathering the Nation and moving us forward in a way which benefits everyone and really keeps our national lead in AI.

Mr. HURD. That's a great way to end our series.

I want to thank our witnesses for appearing before us today. The record is going to remain open for 2 weeks for any member to submit a written opening statement or questions for the record.

And if there's no further business, without objection, the subcommittee stands adjourned.

[Whereupon, at 3:37 p.m., the subcommittee was adjourned.]

