

THE QUANTUM, AI, AND CLOUD LANDSCAPE:
EXAMINING OPPORTUNITIES, VULNERABILITIES,
AND THE FUTURE OF CYBERSECURITY

JOINT HEARING

BEFORE THE

SUBCOMMITTEE ON
CYBERSECURITY AND
INFRASTRUCTURE PROTECTION

AND THE

SUBCOMMITTEE ON
OVERSIGHT, INVESTIGATIONS,
AND ACCOUNTABILITY

OF THE

COMMITTEE ON HOMELAND SECURITY
HOUSE OF REPRESENTATIVES
ONE HUNDRED NINETEENTH CONGRESS

FIRST SESSION

DECEMBER 17, 2025

Serial No. 119-31

Printed for the use of the Committee on Homeland Security



Available via the World Wide Web: <http://www.govinfo.gov>

U.S. GOVERNMENT PUBLISHING OFFICE

63-128 PDF

WASHINGTON : 2026

COMMITTEE ON HOMELAND SECURITY

ANDREW R. GARBARINO, New York, *Chairman*

MICHAEL T. MCCAUL, Texas, <i>Vice Chair</i>	BENNIE G. THOMPSON, Mississippi, <i>Ranking Member</i>
MICHAEL GUEST, Mississippi	ERIC SWALWELL, California
CARLOS A. GIMENEZ, Florida	J. LUIS CORREA, California
AUGUST PFLUGER, Texas	SHRI THANEDAR, Michigan
MARJORIE TAYLOR GREENE, Georgia	SETH MAGAZINER, Rhode Island
TONY GONZALES, Texas	DANIEL S. GOLDMAN, New York
MORGAN LUTTRELL, Texas	DELIA C. RAMIREZ, Illinois
DALE W. STRONG, Alabama	TIMOTHY M. KENNEDY, New York
JOSH BRECHEEN, Oklahoma	LAMONICA MCIVER, New Jersey
ELIJAH CRANE, Arizona	JULIE JOHNSON, Texas, <i>Vice Ranking Member</i>
ANDREW OGLES, Tennessee	PABLO JOSÉ HERNÁNDEZ, Puerto Rico
SHERI BIGGS, South Carolina	NELLIE POI, New Jersey
GABE EVANS, Colorado	JAMES R. WALKINSHAW, Virginia
RYAN MACKENZIE, Pennsylvania	TROY A. CARTER, Louisiana
BRAD KNOTT, North Carolina	AL GREEN, Texas
VINCE FONG, California	
MATT VAN EPPS, Tennessee	

KEIGHLE JOYCE, *Staff Director*
HOPE GOINS, *Minority Staff Director*
SEAN CORCORAN, *Chief Clerk*

SUBCOMMITTEE ON CYBERSECURITY AND INFRASTRUCTURE PROTECTION

ANDREW OGLES, Tennessee, *Chairman*

CARLOS A. GIMENEZ, Florida	ERIC SWALWELL, California, <i>Ranking Member</i>
MORGAN LUTTRELL, Texas	SETH MAGAZINER, Rhode Island
RYAN MACKENZIE, Pennsylvania	LAMONICA MCIVER, New Jersey
VINCE FONG, California	JAMES R. WALKINSHAW, Virginia
ANDREW R. GARBARINO, New York (<i>ex officio</i>)	BENNIE G. THOMPSON, Mississippi (<i>ex officio</i>)

ROLAND HERNANDEZ, *Subcommittee Staff Director*
MOIRA BERGIN, *Minority Subcommittee Staff Director*

SUBCOMMITTEE ON OVERSIGHT, INVESTIGATIONS, AND ACCOUNTABILITY

JOSH BRECHEEN, Oklahoma, *Chairman*

MARJORIE TAYLOR GREENE, Georgia	SHRI THANEDAR, Michigan, <i>Ranking Member</i>
DALE W. STRONG, Alabama	DELIA C. RAMIREZ, Illinois
ANDREW OGLES, Tennessee	TROY A. CARTER, Louisiana
BRAD KNOTT, North Carolina	AL GREEN, Texas
ANDREW R. GARBARINO, New York, (<i>ex officio</i>)	BENNIE G. THOMPSON, Mississippi (<i>ex officio</i>)

GRAYSON WESTMORELAND, *Subcommittee Staff Director*
LISA CANINI, *Minority Subcommittee Staff Director*

CONTENTS

	Page
STATEMENTS	
The Honorable Andrew Ogles, a Representative in Congress From the State of Tennessee, and Chairman, Subcommittee on Cybersecurity and Infrastructure Protection:	
Oral Statement	1
Prepared Statement	3
The Honorable Josh Brecheen, a Representative in Congress From the State of Oklahoma, and Chairman, Subcommittee on Oversight, Investigations, and Accountability:	
Oral Statement	6
Prepared Statement	7
The Honorable Shri Thanedar, a Representative in Congress From the State of Michigan, and Ranking Member, Subcommittee on Oversight, Investigations, and Accountability:	
Oral Statement	4
Prepared Statement	5
The Honorable Bennie G. Thompson, a Representative in Congress From the State of Mississippi, and Ranking Member, Committee on Homeland Security:	
Prepared Statement	8
The Honorable Delia C. Ramirez, a Representative in Congress From the State of Illinois:	
Prepared Statement	8
The Honorable James R. Walkinshaw, a Representative in Congress From the State of Virginia:	
Prepared Statement	9
WITNESSES	
Mr. Logan Graham, Ph.D., Department Head, Frontier Red Team, Anthropic PBC:	
Oral Statement	11
Prepared Statement	12
Mr. Royal Hansen, Vice President, Privacy, Safety, and Security Engineering, Google LLC:	
Oral Statement	17
Prepared Statement	19
Mr. Eddy Zervigon, Chief Executive Officer, Quantum XChange:	
Oral Statement	22
Prepared Statement	24
Mr. Michael Coates, Founding Partner, Seven Hill Ventures:	
Oral Statement	26
Prepared Statement	27
APPENDIX	
Question From Honorable James R. Walkinshaw for Logan Graham	53
Questions From Honorable James R. Walkinshaw for Royal Hansen	53
Question From Honorable James R. Walkinshaw for Eddy Zervigon	54
Questions From Honorable James R. Walkinshaw for Michael Coates	55

THE QUANTUM, AI, AND CLOUD LANDSCAPE: EXAMINING OPPORTUNITIES, VULNERABILITIES, AND THE FUTURE OF CYBERSECURITY

Wednesday, December 17, 2025

U.S. HOUSE OF REPRESENTATIVES,
COMMITTEE ON HOMELAND SECURITY,
SUBCOMMITTEE ON CYBERSECURITY AND
INFRASTRUCTURE PROTECTION, AND THE
SUBCOMMITTEE ON OVERSIGHT,
INVESTIGATIONS, AND ACCOUNTABILITY,
Washington, DC.

The subcommittees met, pursuant to notice, at 10:01 a.m., in room 360, Cannon House Office Building, Hon. Andy Ogles [Chairman of the Cybersecurity and Infrastructure Protection] presiding.

Present from the Subcommittee on Cybersecurity and Infrastructure Protection: Representatives Ogles, Gimenez, Luttrell, Fong, Swalwell, Magaziner, McIver, and Walkinshaw.

Present from the Subcommittee on Oversight, Investigations, and Accountability: Representatives Brecheen, Strong, Ogles, Thanedar, Ramirez, and Carter.

Mr. OGLES. The Committee on Homeland Security, Subcommittee on Cybersecurity and Infrastructure Protection and Subcommittee on Oversight, Investigations, and Accountability will come to order. Without objection, the Chair may declare the committee in recess at any point.

The purpose of this hearing is to examine how rapid advances in artificial intelligence, quantum computing, and cloud technologies are reshaping the cybersecurity landscape in ways that affect both U.S. defensive capabilities and the operational reach of our adversaries. The hearing will also assess how the adoption and governance of AI, cloud infrastructure, and post-quantum security measures are strengthening or, in some cases, exposing U.S. critical infrastructure, Federal systems, and sensitive data, and what steps Government and industry must take to stay ahead of the rapidly-evolving threats.

I now recognize myself for an opening statement.

Good morning and thank you all for being here. I want to begin by thanking Chairman Brecheen and Members of the Subcommittee on Oversight, Investigations, and Accountability for partnering with my subcommittee to hold this hearing. The issues before us today affect national security, economic competitiveness,

and public trust. They deserve attention that reflects their scale and importance.

We are meeting at a time when the technology shaping our digital environment are also shaping the security and strength of the United States. Artificial intelligence, cloud computing, and quantum technologies are now woven into how Federal, State, and local governments operate, how intelligence is collected and analyzed, how critical infrastructure functions, and how American companies compete in a global economy. These technologies offer extraordinary promise, but they also introduce risks that are advancing faster than many of the frameworks and systems designed to manage them.

Artificial intelligence is changing the pace and character of cyber activity. It allows information to be processed at speeds far beyond human capacity and perhaps in some ways even comprehension. It enables automation across complex networks and supports decision making at scale. These capabilities can strengthen cyber defense and improve resilience. However, they can also be exploited to accelerate malicious activities, expand the reach of cyber operations, and make hostile actions more difficult to detect, attribute, and disrupt.

Cloud computing has amplified both opportunity and risk. Cloud platforms have enabled modernization across Government and industry, supporting flexibility, scalability, and innovation. Yet they also consolidate vast amounts of data access and computing power into shared environments, raising the stakes of security configuration and oversight decisions.

Quantum technologies present a longer-term challenge with significant implications. Much of our digital security relies on encryption to protect sensitive communications, verify identities, and secure critical systems. Advances in quantum computing raises serious questions about whether today's encryption methods will remain effective in the future. Our adversaries understand this risk and are already planning, including by collecting encrypted data now with the expectation that it may be accessed later.

The threat environment surrounding these developments is intensifying. The People's Republic of China, PRC, and the Russian Federation, the RF, are investing heavily in advanced computing, automation, and data exploitation as tools of national power. They view artificial intelligence, cloud infrastructure, and emerging technologies as means to gain strategic advantage, conduct sustained cyber and intelligence operations, and operate below the threshold of an open or kinetic conflict.

China, in particular, has pursued a model that tightly integrates government, military, academia, and the private sector. This approach allows innovations developed for commercial purposes to be adapted quickly for state use. In cyber space it supports operations built for scale and persistence, including the use of automated tools to scan networks, identify vulnerabilities, manage stolen credentials, and analyze large volumes of data across many targets simultaneously.

At the same time, these technologies provide the United States with powerful tools to strengthen security and resilience. Artificial intelligence can improve threat detection and response. Cloud com-

puting can enhance reliability and operational flexibility. Advances in quantum research may ultimately yield new security capabilities.

But also there is a downside. The challenge lies in ensuring these benefits are realized without introducing vulnerabilities that adversaries can exploit. The Department of Homeland Security and Cybersecurity and Infrastructure Security Agency, or CISA, play an essential role in this effort. Their work on cloud security practices, artificial intelligence, risk management, and preparation for future changes in encryption help shape how Federal agencies and critical infrastructure operators address emerging threats.

Congress also has an important responsibility. Oversight helps ensure that security keeps pace with adoption that roles—or pace rather—with adoption that roles and responsibilities are clearly defined and that risks are addressed early rather than after they have created serious harm.

This is not about slowing innovation. It is about making sure innovation strengthens the nature rather than exposing it. The decision being made now about how artificial intelligence, cloud computing, and quantum technologies are secured will shape the country's security prosperity for years to come and I would argue, also, our role as the, quite frankly, sole superpower.

I appreciate our witnesses for being here. I look forward to their testimony and the discussion ahead.

[The statement of Chairman Ogles follows:]

STATEMENT OF CHAIRMAN ANDREW OGLES

DECEMBER 17, 2025

Good morning, and thank you all for being here. I want to begin by thanking Chairman Brecheen and the Members of the Subcommittee on Oversight, Investigations, and Accountability for partnering with my subcommittee to hold this hearing. The issues before us today affect national security, economic competitiveness, and public trust, and they deserve attention that reflects their scale and importance.

We are meeting at a time when the technologies shaping our digital environment are also shaping the security and strength of the United States. Artificial intelligence, cloud computing, and quantum technologies are now woven into how Federal, State, and local governments operate, how intelligence is collected and analyzed, how critical infrastructure functions, and how American companies compete in a global economy.

These technologies offer extraordinary promise, but they also introduce risks that are advancing faster than many of the frameworks and systems designed to manage them.

Artificial intelligence is changing the pace and character of cyber activity. It allows information to be processed at speeds far beyond human capacity, enables automation across complex networks, and supports decision making at scale. These capabilities can strengthen cyber defense and improve resilience. However, they can also be exploited to accelerate malicious activity, expand the reach of cyber operations, and make hostile actions more difficult to detect, attribute, and disrupt.

Cloud computing has amplified both opportunity and risk. Cloud platforms have enabled modernization across government and industry, supporting flexibility, scalability, and innovation. Yet, they also consolidate vast amounts of data, access, and computing power into shared environments, raising the stakes of security, configuration, and oversight decisions.

Quantum technologies present a longer-term challenge with significant implications. Much of our digital security relies on encryption to protect sensitive communications, verify identities, and secure critical systems. Advances in quantum computing raise serious questions about whether today's encryption methods will remain effective in the future. Our adversaries understand this risk and are already planning for it, including by collecting encrypted data now with the expectation that it may be accessed later.

The threat environment surrounding these developments is intensifying.

The People's Republic of China and the Russian Federation are investing heavily in advanced computing, automation, and data exploitation as tools of national power. They view artificial intelligence, cloud infrastructure, and emerging technologies as means to gain strategic advantage, conduct sustained cyber and intelligence operations, and operate below the threshold of open conflict.

China, in particular, has pursued a model that tightly integrates government, military, academia, and the private sector. This approach allows innovations developed for commercial purposes to be adapted quickly for State use. In cyber space, it supports operations built for scale and persistence, including the use of automated tools to scan networks, identify vulnerabilities, manage stolen credentials, and analyze large volumes of data across many targets simultaneously.

At the same time, these technologies provide the United States with powerful tools to strengthen security and resilience. Artificial intelligence can improve threat detection and response. Cloud computing can enhance reliability and operational flexibility. Advances in quantum research may ultimately yield new security capabilities. The challenge lies in ensuring these benefits are realized without introducing vulnerabilities that adversaries can exploit.

The Department of Homeland Security and the Cybersecurity and Infrastructure Security Agency play an essential role in this effort. Their work on cloud security practices, artificial intelligence risk management, and preparation for future changes in encryption helps shape how Federal agencies and critical infrastructure operators address emerging threats.

Congress also has an important responsibility. Oversight helps ensure that security keeps pace with adoption, that roles and responsibilities are clearly defined, and that risks are addressed early rather than after serious harm has occurred. This is not about slowing innovation. It is about making sure innovation strengthens the Nation rather than exposing it.

The decisions being made now about how artificial intelligence, cloud computing, and quantum technologies are secured will shape the country's security and prosperity for years to come.

Mr. OGLES. I now recognize the Ranking Member for the Subcommittee on Oversight, Investigations, and Accountability, the gentleman from Michigan, Mr. Thanedar, for his opening statement.

Mr. THANEDAR. Thank you, Chairman Ogles. Appreciate this hearing. Good morning to all of our witnesses. I look forward to hearing your thoughts.

For two decades, hostile nations have conducted increasingly sophisticated cyber attacks against the United States. These attacks have been used to spy, steal intellectual property, cripple critical infrastructure, and demand ransom payments. China, Russia, Iran, North Korea are aggressively using advanced cyber capabilities to threaten our national security and economic prosperity. China is both the most active and persistent cyber threat and is also the only country with both the desire and the ability to reshape the world order, which is why it is extremely shocking that President Trump recently agreed to allow Nvidia to sell advanced artificial intelligence chips to China. Really shocking.

Let's just see some background information here. Why did this decision the President made? The President was quick to sell out America's security after Nvidia's CEO attended a \$1 million per plate dinner at Mar-a-Lago and donated to Trump's White House ballroom. So much for America First.

Trump's own Department of Justice has warned that China is seeking to become the AI leader by 2030 and plans to use AI chips to modernize its military, design and test weapons of mass destruction, and deploy advanced surveillance tools. We should be disrupting and dismantling threat actors whose actions threaten our national interest, not enabling them.

The rapid development of emerging technologies, including advanced AI and quantum computing, enables and enhances security risks. These advanced technologies not only accelerate the cyber abilities of countries such as China, but they also make it easier for countries that are not well-resourced and enable a growing threat from organized criminal groups.

Over the past year, cyber attacks have become faster, more widespread, and harder to detect. As AI-assisted cyber attacks hit harder and faster, it is critical that Congress extends CISA 2015, the Cybersecurity Information Sharing Act of 2015. CISA 2015 provides privacy and liability protection to companies to encourage them to share data about cyber vulnerabilities and threats. These protections are necessary to fully understand the risk and facilitate collaboration between the Federal Government and the private sector.

Unfortunately, CISA 2015 expires next month. A 10-year extension is the best reauthorization strategy that will also provide the private sector with assurances while eliminating the risk of this authority lapsing.

I look forward to hearing from our witnesses how else we can best defend against cyber attacks that are leveraging powerful emerging technologies.

Thank you and I yield back, Mr. Chair.

[The statement of Ranking Member Thanedar follows:]

STATEMENT OF RANKING MEMBER SHRI THANEDAR

DECEMBER 17, 2025

For two decades, hostile nations have conducted increasingly sophisticated cyber attacks against the United States. These attacks have been used to spy, steal intellectual property, cripple critical infrastructure, and demand ransom payments.

China, Russia, Iran, and North Korea are aggressively using advanced cyber capabilities to threaten our national security and economic prosperity. China is both the most active and persistent cyber threat and is also the only country with both the desire and ability to reshape the world order. Which it is why it shocking that President Trump recently agreed to allow Nvidia to sell advanced artificial intelligence chips to China.

The President was quick to sell out America's security after Nvidia's CEO attended a \$1 million-per-plate dinner at Mar-a-Lago and donated to Trump's White House ballroom boondoggle. So much for "America First"!

Trump's own Department of Justice has warned that China is seeking to become the AI leader by 2030 and plans to use AI chips to modernize its military, design and test weapons of mass destruction, and deploy advanced surveillance tools. We should be disrupting and dismantling threat actors whose actions threaten our national interests, not enabling them.

The rapid development of emerging technologies, including advanced AI and quantum computing, enables and enhances security risks. These advanced technologies not only accelerate the cyber abilities of countries such as China, but they also make it easier for countries that are not well-resourced and enable a growing threat from organized criminal groups.

Over the past year, cyber attacks have become faster, more wide-spread, and harder to detect. As AI assisted cyber attacks hit harder and faster, it is critical that Congress extend CISA 2015—the Cybersecurity Information Sharing Act of 2015. CISA 2015 provides privacy and liability protections to companies to encourage them to share data about cyber vulnerabilities and threats. These protections are necessary to fully understand the risks and facilitate collaboration between the Federal Government and the private sector.

Unfortunately, CISA 2015 expires next month. A 10-year extension is the best reauthorization strategy that will also provide the private sector with assurances while eliminating the risk of this authority lapsing. I look forward to hearing from

our witnesses how else we can best defend against cyber attacks that are leveraging powerful emerging technologies.

Mr. OGLES. Thank you, Ranking Member Thanedar, and I look forward to following up on your insightful comments.

I now recognize the Chairman for the Subcommittee on Oversight, Investigations, and Accountability, the gentleman from Oklahoma, Mr. Brecheen, for his opening statement.

Mr. BRECHEEN. Thank you, Chairman Ogles. Good morning. Thank you to our witnesses. Very complex subject. Many of us in all vulnerability feel really unqualified to be in these discussions. Grateful we are going to have some expertise to drive into the massive amount of vulnerabilities that AI is presenting on our cyber front. As Chair of the Subcommittee on Oversight, Investigations, and Accountability, I am looking forward to partnering with the Subcommittee on Infrastructure Protection to focus on this topic, explore ways that Congress can assist the Department of Homeland Security in countering this new threat.

This integration of AI into cyber attacks should concern every American. The recent cyber attack leveraging Anthropic's AI infrastructure showed that complex attack campaigns can now be conducted with little to no human interaction, at speeds faster than a human could replicate. We have all seen how AI can easily streamline tasks that would otherwise be very labor intensive, both in business and everyday life. Now that an attack like this has successfully taken place, we can expect to see more events like this in the future. The proof of concept is there, and even if U.S.-based AI companies can put safeguards against using their models for such attacks, these actors will find other ways to access this technology.

China is our most significant cyber threat actor and it continues to search for tactics to infiltrate critical U.S. systems and prioritize the development of advanced computing technology and AI that supports its economic and strategic goals. Cyber espionage has been a key part of their plan, China's plan ongoing campaign of stealing intellectual property. This is decades-old and they now have new tools, and this will fuel rapid technological advancement at the expense of American innovators.

As this committee has highlighted over the years, cyber actors linked to China pose a threat on an unprecedented scale targeting U.S. companies, critical infrastructure, and the Federal Government. As technologies like AI continue to advance at such speeds, we have to be vigilant strategic in protecting intellectual property and our national security. From an oversight perspective, we need to make sure that Federal civilian agencies are taking the proactive steps needed to protect their networks against intrusion. Technology doesn't advance on the Government's time line and we can't afford to have cybersecurity practices moving at such speeds absent Government interdiction. That path leaves us reacting to security failures instead of proactively confronting today's threats.

This is an area where Federal Government can partner with and learn from the private sector to implement best practices and incorporate needed technology. The Federal Government needs to be better at sharing information on cyber threats between Federal agencies and with private stakeholders in a timelier manner. I hope to learn in today's hearings how Congress can empower the

Department of Homeland Security and its sub-agencies to counter this threat and ensure safety integrity of U.S.-based infrastructure.

I want to thank again our panel of witnesses for joining us to discuss today to discuss the cyber attack, implementation of that. Congress and American people need to consider how we can work with you all in your expertise to safeguard our critical infrastructure.

With that, I want to yield back to Chairman Ogles.
[The statement of Chairman Brecheen follows:]

STATEMENT OF CHAIRMAN JOSH BRECHEEN

DECEMBER 17, 2025

Thank you, Chairman Ogles. Good morning and thank you for joining us today to discuss the highly complex and important issue of artificial intelligence's role in carrying out cyber attacks.

As Chair of the subcommittee on Oversight, Investigations, and Accountability, I am looking forward to partnering with our Subcommittee on Cybersecurity and Infrastructure Protection to focus on this topic and explore ways Congress can assist the Department of Homeland Security in countering this new threat.

The integration of AI into cyber attacks should concern all Americans.

The recent cyber attack leveraging Anthropic's AI infrastructure showed that complex attack campaigns can now be conducted with little-to-no human intervention at speeds faster than any human could replicate.

We've all seen how AI can easily streamline tasks that would otherwise be labor-intensive, both in business and in everyday life.

However, now that an attack like this has successfully taken place, I think we can expect to see more events like this in the future.

The proof of concept is there. And even if U.S.-based AI companies can put safeguards against using their models for cyber attacks, cyber threat actors will find other ways to access this technology.

China, our most significant cyber threat actor, continues to search for new tactics to infiltrate critical U.S. systems, and is prioritizing the development of advanced computing technology and AI that supports its economic and strategic goals.

Cyber espionage has been a key part of China's on-going campaign of stealing intellectual property to fuel rapid technological advancement at the expense of American innovators.

As this committee has highlighted over the years, cyber actors linked to China pose a threat on an unprecedented scale targeting U.S. companies, critical infrastructure, and the Federal Government.

As technologies like AI continue to advance at rapid speeds, we must be vigilant and strategic in protecting our intellectual property and national security.

From an oversight perspective, we need to make sure that Federal civilian agencies are taking the proactive steps needed to protect sensitive networks against intrusions.

Technology doesn't advance on the Government's time line; we can't afford to have Federal cybersecurity practices move at the speed of government.

That path leaves us reacting to security failures instead of proactively confronting today's evolving threats.

This is an area where the Federal Government can partner with, and learn from, the private sector to implement best practices and incorporate modern technology.

Additionally, the Federal Government needs to be better at sharing information on cyber threats between Federal agencies and with private stakeholders, in a timelier manner.

I hope to learn in today's hearing ways that Congress can empower the Department of Homeland Security, and its subagencies, to counter this threat and ensure the safety and integrity of U.S.-based cyber infrastructure.

I want to thank our panel of witnesses for joining us today to discuss this latest cyber attack and the implications that Congress and the American people need to consider as we think about how to protect critical networks in the age of AI.

Mr. OGLES. Thank you, Chairman Brecheen, and just echo your sentiments.

Other Members of the committee, you are reminded that you can submit for the record an opening statement.

[The statements of Ranking Member Thompson, Honorable Ramirez and Honorable Walkinshaw follow:]

STATEMENT OF RANKING MEMBER BENNIE G. THOMPSON

DECEMBER 17, 2025

With cybersecurity threats constantly evolving, it is essential that we assess how to stay ahead of our adversaries by both defending against new technological threats and by developing and deploying the best tools to defend our networks.

Anthropic's recent report on the use of AI by Chinese state-sponsored actors demonstrates just how rapidly changes in technology can impact cybersecurity. Since ChatGPT's launch just 3 years ago, we have already seen large language models significantly change how hackers carry out cyber campaigns.

Today's hearing will give the subcommittee an opportunity to hear from the private sector on how a range of technological innovations are impacting cybersecurity today and how the Federal Government can better prepare for tomorrow's threats. While I appreciate the strong bipartisan interest in this topic, I worry that many of the Trump administration's actions are moving us in the wrong direction by hamstringing both the public and private-sector security efforts. CISA has lost hundreds of employees this year, and during the shutdown, the administration attempted to illegally fire CISA's stakeholder engagement staff—the very staff who carry out the public-private collaboration we all agree is necessary in cybersecurity.

Across the Federal Government, the administration's war against Federal employees has reduced technological expertise and done long-term damage to the Federal Government's ability to recruit and retain technology experts. I cannot imagine the Chinese government would try to force out their AI or quantum experts, yet that is exactly what we have seen the Trump administration do here. Such actions make us less safe and put us at a competitive disadvantage.

At the same time, the Trump administration has implemented anti-immigrant policies that have made it harder for high-skilled immigrants to move the United States, while harassing and profiling immigrants already living here. The United States will never be able to compete with China on the size of our overall workforce.

But, our ability to attract the best and the brightest from around the world has always given us an advantage, as we have seen from the many immigrants who have founded and led cutting-edge technology companies. If we close the door to immigrants, our national security will suffer. As we in Congress assess how to strengthen our cybersecurity, we must increase our oversight over CISA and other Federal agencies to better understand how they are combatting new threats with their current staffing and resources and how recent policy decisions have impacted our security posture. I hope that we will have CISA officials before the committee soon we can ask them these important questions.

Additionally, we must fulfill our obligations to maintain and grow our Nation's cybersecurity capacities by passing a long-term reauthorization of the Cybersecurity Information Sharing Act of 2015, while adequately funding research and development in novel technologies and security. As we consider oversight and legislative activities next year, I am confident the witnesses' testimony today will help inform our efforts.

STATEMENT OF HONORABLE DELIA C. RAMIREZ

DECEMBER 17, 2025

Thank you, Chair and Ranking Member, for holding today's hearing, and to our witnesses for joining us.

It's really hard to talk about the "opportunities" around AI, quantum, and cloud computing to reduce the risk of cyber threats, when the Department of Homeland Security (DHS) is using similar technologies and its own private partnerships to threaten our communities.

DHS is violating our rights with AI, data monitoring, and surveillance technologies they've purchased with taxpayer dollars:

1. DHS kept Chicago Police records in direct violation of domestic espionage rules designed to prevent domestic intelligence operations from targeting legal U.S. residents.

2. In Chicago, DHS is using facial recognition technology to target immigrants while removing policies intended to restrain their use from their website.

a. In Chicago and the State of Illinois, Clearview AI is banned from doing business with police agencies because of a lawsuit that alleged they violated a landmark State law protecting our personal information. This is the same company DHS now has a \$9.2 million dollar contract with.

3. DHS has also used AI technology to do “AI assisted reviews” of social media in what is described as a surveillance program on a scale that was never possible before, and has the potential to have a chilling effect on free speech on a never-before-seen scale.

DHS’s use of technology, data, and AI to surveil our communities and suppress dissent is deeply alarming. But it is unsurprising, given the lawlessness demonstrated by Trump, Secretary Noem, and DHS leadership. That is why it is critical that we do not ignore the opinion and expertise of the privacy, technology, and civil rights experts who are calling out the threat that DHS’s unregulated, unaccountable, unlawful use of technology poses to data protections, privacy, and civil rights.

Whether it’s scanning social media accounts or tracking people’s movements, it is evident that AI is being used to target communities who dissent and to execute Trump’s racist and xenophobic mass deportation campaign. It is critical that meaningful restrictions be put in place. That requires limitations on government use, specifically, but also requires AI developers to limit how their technology is used by consumers.

It’s laughable that the Republicans—the party of small government—are totally comfortable being the party of big brother.

If you ask Republicans:

1. Want to use the power of Government to end hunger? No.

2. Want to use it to address climate change? No.

3. Want to use it to end homelessness? No.

If you ask Republicans, if they want to use it to strike fear into the hearts of your people, chill dissent, and undermine the foundations of liberty and democracy? Why, yes. Yes, let’s do that.

STATEMENT OF HONORABLE JAMES R. WALKINSHAW

WEDNESDAY, DECEMBER 17, 2025

A highly-skilled workforce, combined with the adoption of cutting-edge technologies, should be the foundation of our Nation’s efforts to remain the global leader in emerging technology and to counter cyber threats to our national security. Unfortunately, the Trump administration has purged much of its technical expertise through Department of Government Efficiency (DOGE) “reductions in force” (RIFs) and its Deferred Resignation Program (DRP). Entire units such as 18F were entirely eliminated. Engineers, data scientists, and designers from the U.S. Digital Service have been laid off. Hundreds quit because of the organizational chaos created by President Trump and DOGE. Artificial Intelligence (AI) experts brought in under the National AI Talent surge were pushed out. Just as we need AI and cyber talent the most, the Trump administration has fired and driven them out.

The Trump administration is now promoting its new “Tech Force” program, which would bring private-sector tech talent into Government for short-term stints, as a magical solution to Federal modernization challenges. This administration’s assault on the Federal workforce has made it almost impossible to recruit the highly-skilled and highly-knowledgeable people that we need to make our Government work and to counter cyber threats we are facing today. Short-term hiring initiatives like “Tech Force” will not repair the lasting damage this administration has inflicted on the Federal Government’s ability to recruit and retain technical talent needed to meet evolving national security threats from malign actors.

The United States must prioritize maintaining a sophisticated Federal workforce to ensure we remain positioned to deter cyber attacks.

Mr. OGLES. I am pleased to have a distinguished panel of witnesses before us today on this critical topic. Pursuant to committee rule VII(C), I ask that our witnesses please rise and raise their right hands.

[Witnesses sworn.]

Mr. OGLES. Let the record reflect that the witnesses have answered in the affirmative. Thank you and please be seated.

I would like to now formally introduce our witnesses.

Dr. Logan Graham serves as the department head of the Frontier Red Team at Anthropic, where he leads efforts to evaluate the behavior and potential misuse of advanced AI systems as model capabilities continue to scale. His work focuses on identifying national security risks posed by Frontier AI, including its potential use in cyber espionage and offensive cyber operations, as well as developing safeguards to detect and disrupt malicious activity.

Prior to joining Anthropic, Dr. Graham held roles at Google X and Babylon Health. He also previously served as special advisor to the Prime Minister of the United Kingdom, contributing to national science and technology policy, and the development of the United Kingdom's AI strategy. Dr. Graham earned his undergraduate degree in economics from the University of British Columbia and completed his Ph.D. in engineering science at the University of Oxford where he was a Rhodes Scholar. Thank you, sir.

Mr. Royal Hansen is vice president for privacy, safety, security engineering at Google, where he leads the engineering team research responsible for securing Google's global technical infrastructure and protecting billions of users world-wide. Prior to joining Google, Mr. Hansen held senior security leadership roles in the financial services sector, including at American Express, Goldman Sachs, Morgan Stanley, and Fidelity Investments. Mr. Hansen holds a bachelor of arts in computer science from Yale University. Thank you, sir.

Mr. Eddy Zervigon is the chief executive officer of Quantum XChange. Under his leadership, Quantum XChange works with Government and private-sector partners to prepare critical systems for emerging cyber- and quantum-enabled threats. Mr. Zervigon brings extensive experience in corporate leadership, operations, and restructuring, including prior service as a managing director in the Principal Investments Group at Morgan Stanley, where he oversaw technology and infrastructure investments across the United States and Latin America. He holds a bachelor's degree in accounting and a master's degree in taxation from Florida International University and a master of business administration from Dartmouth. Thank you, sir.

Mr. Michael Coates is the founding partner of Seven Hill Ventures, an early-stage venture firm focused exclusively on cybersecurity investment, addressing enterprise operational and national security challenges. He brings more than two decades of experience securing large-scale digital platforms and advising organizations on cyber risk.

Mr. Coates previously served as the chief information officer at Twitter and also led security efforts at Mozilla. Mr. Coates holds a bachelor of science in computer science from the University of Illinois Urbana-Champaign and a master of science in computer information and network security from DePaul University. I thank each of our—thank you, sir.

I thank each of our distinguished witnesses for being here today.

This is a topic that, you know, a year-and-a-half ago was somewhat of a niche for laypersons, but for you experts, obviously, clearly recognize that this was going to be, quite frankly, the next arms race, threat battlefield as we go forward. So what you are doing

today here before Congress means more than I think we can possibly comprehend as we begin this discussion and, quite frankly, dive into the emergence of this technology.

With that, I now recognize Dr. Graham for 5 minutes to summarize his opening statement.

**STATEMENT OF LOGAN GRAHAM, PH.D., DEPARTMENT HEAD,
FRONTIER RED TEAM, ANTHROPIC PBC**

Mr. GRAHAM. Chair Ogles and Brecheen, Ranking Member Thanedar, Members of the committee, thank you for the opportunity to testify today.

Anthropic is a leading Frontier AI model developer working to build reliable, interpretable, and steerable artificial intelligence. Our flagship AI assistant, Claude, serves millions of Americans and trusted partners worldwide, from Fortune 500 companies and U.S. Government agencies to small businesses and cutting-edge startups, and consumers, enhancing productivity on tasks including software engineering, data analysis, and scientific research.

At Anthropic, I lead the Frontier Red Team. Our job is to build an early warning system for advanced risks from AI so that we can mitigate them and to help the world prepare as far in advance as possible. Transparency is a fundamental value for Anthropic and we believe it should be an industry standard. That is why we published a report about how in mid-September 2025 anthropic detected suspicious activity that our investigation determined to be a largely autonomous, sophisticated cyber espionage campaign conducted by a group sponsored by the Chinese Communist Party.

To be clear, Claude's code was not compromised, nor were Anthropic's labs infiltrated. Instead, this group maliciously misused Claude to automate large portions of cyber attacks against their targets. We estimate their use of the model allowed them to automate approximately 80 to 90 percent of the work that previously required humans to do. This is a significant increase in the speed and scale of operations compared to traditional methods.

Further, this group invested significant resources and used our sophisticated network—used their sophisticated network infrastructure in order to circumvent our safeguards and detection mechanisms prior to being detected. They then deceived the model into believing the tasks were ethical cybersecurity tests. The campaign consisted of a few distinct phases. First, a human operator provided targets to Claude, directing it to conduct autonomous reconnaissance against them in parallel. Second, acting on the human operator's direction, Claude leveraged third-party software tools to search for vulnerabilities in these systems. The third and final step was to task Claude to exploit these vulnerabilities and extract sensitive information from the targets, which was only successful in a handful of cases.

We detected this campaign. Within 2 weeks, the attackers first confirmed offensive activity, triggering a swift response, including account bans, strengthening our safeguards, entity notifications, authority coordination, and indicator sharing with partners.

We have reached an inflection point in cybersecurity. It is now clear that sophisticated actors will attempt to use AI models to enable cyber attacks at unprecedented scale. This threat is not

unique to Claude and affects all AI models. That is why we've been open and transparent about this incident and one of the reasons why I'm grateful to you that you are holding this hearing today. Industry and Government must collaborate to prevent this misuse and enable cyber defenders to prepare.

To address these risks there are at least 3 things that should be done immediately. First, there needs to be rapid testing of models for national security capabilities. Government-led evaluations, like those conducted by NIST's Center for AI Standards and Innovation, give us visibility into model capabilities and security. Codifying and expanding this process is critical.

Second, there must be robust threat intelligence sharing. Frontier AI labs and the U.S. Government need stronger channels to share indicators of misuse as exists in critical infrastructure sectors.

Third, and finally, industry should invest in empowering our cyber defenders. We must make models useful for defenders and get them into their hands. Anthropic is improving its models for cyber defenders and building tools, for example, that can patch vulnerabilities.

We cannot lose sight of the strategic picture. The United States and its allies must maintain leadership in AI. The Trump administration has taken important steps to advance U.S. AI leadership, including accelerating the build-out of AI infrastructure, promoting Federal adoption, and strengthening security testing and coordination. We strongly support these efforts.

Equally critical is maintaining the United States' advantage in computing power, the single most important input into developing powerful AI models. The United States currently has a significant edge over the CCP in access to advanced chips. But if advanced compute flows to the CCP, its national champions could train models that exceed U.S. frontier cyber capabilities. Attacks from these models will be much more difficult to detect and deter.

We are in a race against threat actors who will stop at nothing to misuse AI for cyber attacks. Our response must be urgent, coordinated, and focused on securing systems faster than they can be attacked.

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

[The prepared statement of Mr. Graham follows:]

PREPARED STATEMENT OF LOGAN GRAHAM

DECEMBER 17, 2025

Chair Ogles, Chair Brecheen, Ranking Member Swalwell, Ranking Member Thanedar, and Members of the committee, thank you for the privilege and opportunity to testify today.

Anthropic is a leading frontier AI model developer working to build reliable, interpretable, and steerable artificial intelligence (AI) systems. Anthropic has become the fourth-most valuable private company in the world.¹ Our flagship AI assistant, Claude, serves millions of Americans and trusted partners worldwide, from Fortune 500 companies and U.S. Government agencies to small businesses, cutting-edge

¹Yuliya Chernova, "Anthropic Valuation Hits \$183 Billion in New \$13 Billion Funding Round," *The Wall Street Journal*, Sept. 2, 2025, www.wsj.com/articles/anthropic-valuation-hits-183-billion-in-new-13-billion-funding-round-6212f3ed.

startups and consumers, enhancing productivity on sophisticated tasks including software development, data analysis, and scientific research.

We believe these AI models could become extremely powerful very soon. We think that by late 2026 or early 2027, it may be possible to have “a country of geniuses in a data center.” America is in an excellent position to lead its development, and we must preserve this advantage.

The benefits of powerful AI will be immense. We see it enabling pioneering cancer research, supporting discoveries in material science, and providing health care support where it’s most needed. AI is now unlocking large productivity increases for the world’s largest businesses, as well as small and nimble start-ups. Anthropic is committed to making these benefits available to the world while safely and securely stewarding the development of powerful AI.

I lead Anthropic’s Frontier Red Team, an internal research team that studies the capabilities of frontier AI models. Our work generates insights that enable rapid, responsible AI development and inform policy on frontier AI capabilities and risks. The team focuses its evaluations in three critical domains: cybersecurity capabilities, biosecurity risks, and increasing autonomy in AI models. We primarily evaluate Anthropic’s Claude series of frontier models, but in some circumstances evaluate models from other AI developers. Our work shows that AI models are rapidly becoming more capable in areas like cybersecurity—capabilities that, in the right hands, can dramatically strengthen our U.S. and allied national security.

My team has been tracking cybersecurity capabilities of AI models since late 2022. We were among the first in the world to study the dramatic cybersecurity implications of a world where models match or exceed humans in these capabilities. We have allocated significant resources to studying and experimenting on model cybersecurity capabilities. In essence, this amounts to testing AI models’ capabilities by giving them the same hacking tasks you might give to a human. In those tests, we have seen a very consistent trend: models have shown rapid progress on cybersecurity challenges. Two years ago, models were largely unable to complete most basic cybersecurity tasks; last year, they began to do so reliably; and this year, they have begun outcompeting humans in some head-to-head competitions.

We are confident that now is the moment to act. Anthropic is determined to support defenders, and we believe that other model developers, cybersecurity companies and researchers, and the United States Government all have important roles to play. We must also take whatever steps are necessary to ensure America maintains its lead in developing powerful AI, including restricting our adversaries’ access to advanced AI chips and the tools needed to manufacture them. These types of controls are vital to our national security and economic competitiveness.

Today, I will discuss how Anthropic discovered, disrupted, and publicly disclosed what we believe is the first documented case of a successful, highly autonomous cyber espionage campaign that relied on the misuse of AI models. We assess with high confidence that this campaign was conducted by a highly-sophisticated Chinese Communist Party (CCP)-sponsored group. This cyber espionage campaign demonstrates that a sophisticated, well-resourced threat actor—one willing to go to great lengths to circumvent AI model safeguards and deceive the AI model about its true intentions—can now extract meaningful operational value from frontier AI models.

We believe this is the first indicator of a future where, despite strong safeguards, AI models may enable threat actors to conduct an unprecedented scale of cyber attacks, and that these cyber attacks may become increasingly sophisticated in their nature and scale.

AI-DRIVEN CYBER ESPIONAGE CAMPAIGN SPONSORED BY THE CCP

In mid-September 2025, Anthropic detected a sophisticated cyber espionage operation where malicious actors abused our model, Claude, in violation of Anthropic’s Acceptable Use Policy.² While we have safeguards in place designed to detect and prevent this kind of malicious activity, in this case we were confronted with a sophisticated and well-resourced effort to circumvent those defenses and manipulate Claude into complying with the attackers’ instructions.

A CCP-sponsored group misused Claude to automate a substantial part of the process of conducting the attacks. Based on our investigation, we believe the attacks targeted roughly 30 entities, with the goal of finding and extracting valuable information from these entities. While a majority of these infiltration attempts failed, a small number were successful. Upon detecting this attack, we launched an investigation, disrupted the campaign, implemented new mitigations to prevent similar

²“Usage Policy.” Anthropic, Sept. 15, 2025, <https://www.anthropic.com/legal/aup>.

activity, coordinated with the authorities, notified affected entities, and shared technical indicators with our partners to mitigate similar campaigns.

We believe that this group’s abuse of Claude was able to substantially increase the speed and scale of the attack. Importantly, however, our takeaway is that this is not a story just about Claude, nor about what the attack was able to accomplish.

This challenge is not unique to Anthropic—every frontier model developer will face increasingly sophisticated attempts by threat actors to circumvent safeguards and misuse their models. What we observed here is one data point on a trendline. As models become more capable, we expect a wider swath of threat actors will continue to seek ways to misuse models for malicious ends. That is why the entire industry, along with government partners, must continue to strengthen our defenses.

DETAILS OF THE CCP-BACKED CYBER ESPIONAGE CAMPAIGN

The attackers developed a framework designed to execute components of their cyber espionage campaign in a way that relied on human input at a few key points but which was able to misuse Claude Code (a popular product of ours that enables Claude to autonomously write and execute code) and open standard Model Context Protocol (MCP) tools to execute many components of the cyber espionage campaign with a substantial degree of autonomy.³ Using this combination of tools, the attackers circumvented our safeguards and deceived the model about the true nature of the tasks they were directing Claude to complete.

The campaign consisted of distinct phases. At first, a human operator input a target—for example, an entity, or an entity’s network—to Claude. The framework’s orchestration engine would then task Claude to autonomously conduct reconnaissance against multiple targets in parallel. Approximately 30 systems from foreign governments and global companies were targeted, consistent with the threat actor’s instructions. Upon completion, Claude delivered results to the operators for review and to determine the next step.

Next, acting on the threat actor’s direction, Claude leveraged third-party software tools to search for vulnerabilities in these systems. Claude looked for “weak spots” in the target’s infrastructure that could be exploited for the operators to gain unauthorized access to these systems. Many of these software tools were the same open-source software tools used by legitimate defensive actors.

The next and final step was to attempt to exploit any discovered vulnerabilities using third-party tools and to then find and extract sensitive information. This was only successful in a handful of cases, but required similar abilities to scan for systems containing valuable information, identify and exploit vulnerabilities, and exfiltrate the information. It also involved “moving laterally” within the system to establish access to new areas of the target’s system. At the threat actor’s direction, Claude queried databases, extracted information, parsed results to identify proprietary information, and categorized findings by intelligence value to the human operator. Claude then produced a summary report for the human operators to review.

This attack demonstrated that current frontier AI models are capable of uplifting dedicated, sophisticated groups.⁴ Our preliminary estimate is that the threat actor was able to leverage Claude to perform the work of a 10-person team managed by one human operator. For example, we observed that approximately 80 to 90 percent of the CCP-backed campaign tasks were automated by Claude, whereas the remaining 10 to 20 percent were tasks where the human operators reviewed Claude’s outputs and directed the models.

There were critical limitations in the campaign. First, the models frequently hallucinated. Hallucinations are when models essentially “make up” incorrect information—in this case, false credentials, or that it had succeeded when in reality it had not. This means human operators have to spend more time carefully validating all claimed results, limiting overall operational effectiveness. Second, the attack still fundamentally required a human operator at various decision points to progress. That is, the models still requested approval to progress from reconnaissance to active exploitation, authorize use of harvested credentials, and to make final decisions about data exfiltration. Last, the campaign did not produce fundamentally novel attack techniques unknown to security practitioners. Rather, it applied existing methods to identify and exploit vulnerabilities in software systems at scale.

³“Introducing the Model Context Protocol.” Anthropic, Nov. 24, 2024, <https://www.anthropic.com/news/model-context-protocol>.

⁴“Uplift” is the term we use to estimate how much individuals are able to benefit from using models compared to if they had tried to accomplish the same outcome without using models.

ANTHROPIC’S WORK TO DISRUPT THE CCP-BACKED ESPIONAGE CAMPAIGN

Anthropic detected this CCP-backed campaign within 2 weeks of the attackers’ first confirmed offensive activity. Anthropic maintains multiple systems designed to detect suspicious activity, including cyber classifiers and what are known as YARA rules in the security industry.⁵ In this case, one of these systems triggered an immediate human investigation. Over the following 10 days, we banned the associated accounts, implemented detection mechanisms for similar behavior, notified affected entities, and coordinated with authorities to gather actionable intelligence. We also collected the technical indicators of these attacks, and took steps to share these with partners, including other frontier labs, with whom we have threat-sharing agreements, so that they could identify and mitigate similar campaigns.

We assessed with high confidence that the threat actor was affiliated with the CCP because of technical evidence from the sophisticated obfuscation infrastructure that enabled the threat actor to access Claude accounts and evade detection. In addition, the targeted entities aligned with known targets of the CCP; and the operators exhibited behavior consistent with this conclusion, including following the Chinese workday—including observing lunch breaks—and observing Chinese national holidays.

The threat actor went to great lengths to obfuscate their work, conceal their intentions from Claude, or evade our safeguards. First, the actor “jailbroke” our models by, in some instances, deceiving the model, falsely stating they were conducting ethical defensive cybersecurity testing. Then, having convinced the models to comply, the attackers created a sophisticated network of many accounts, which all used separate instances of the model to perform subcomponents of the attacks on different targets. Separating work in this way frequently makes the subcomponents seem benign, but when put together, form a pattern of malicious behavior. They routed their actions through an obfuscated network they controlled.

ANTHROPIC IS CONTINUING TO SECURE ITS MODELS IN RESPONSE TO THIS CAMPAIGN

During and after the campaign, we instituted new mitigations to better prevent this kind of misuse of Anthropic models. We expanded our detection mechanisms to better cover novel threats such as this campaign—including by improving our cyber-focused classifiers. We are also prototyping early detection systems specifically targeted at autonomous cyber attacks, and researching new techniques for investigating and mitigating large-scale distributed operations.

Importantly, because all AI models are susceptible to this type of misuse, we shared and continue to share the results of our investigation with frontier labs. Defensive actors world-wide need to prepare for and defend against these new threats.

WHAT INDUSTRY AND GOVERNMENT SHOULD DO

As model capabilities advance, AI developers have to get better at understanding risks, preventing misuse, and ensuring that models can be used by defenders. This is a shared challenge on which industry and government should work together. While the threat actors likely leveraged Claude for this campaign due to its advanced coding and agentic capabilities, many models available today could soon be able to conduct such an attack. It is therefore critical that industry, Government, and researchers work together to evaluate model capabilities, rapidly secure critical infrastructure, and develop better methods to restrict malicious use.

Predeployment Testing and Transparency for National Security Capabilities

The United States should continue to be the best and fastest at evaluating model capabilities, deploying models, and learning from these deployments. Government-led evaluations remain critical, as the intelligence community and agencies like the Department of Energy possess unique expertise to evaluate how adversaries could exploit AI models.

The Frontier Red Team has an on-going partnership with the U.S. Government that enables risk mitigation and provides strategic national security insights. One major part of this is our collaboration with the U.S. Center for AI Standards and Innovation (CAISI) in the Department of Commerce. Through voluntary agreements, the CAISI conducts rapid predeployment testing of our Claude models that gives the Government visibility into AI model capabilities, provides us with critical information about our models’ national security implications, and allows us to launch our commercial models more rapidly and with enhanced confidence about

⁵“Using YARA For Malware Detection.” NCCIC, https://www.cisa.gov/sites/default/files/FactSheets/NCCIC%20ICS_FactSheet_YARA_S508C.pdf.

their reliability. Because of the sensitive nature of cybersecurity information, the CAISI and the U.S. Government in general are in an advantageous position to evaluate model capabilities and understand capability trajectories better than anyone in the world. Codifying the CAISI can ensure the Government can test and evaluate models for these capabilities, in partnership with the U.S. national security community.

In conjunction with Government testing, transparency standards play a crucial role in achieving secure AI development. This is why Anthropic published a transparency framework to inform light-touch guardrails that encourage the largest AI developers to follow secure practices—disclosing how they assess and mitigate national security risks, their testing procedures, and results.⁶ This transparency approach would establish industry best practices for safety and set a baseline for secure model training, ensuring developers meet basic accountability standards while enabling public visibility into development without impeding innovation.

Threat Intelligence Sharing

Additionally, the U.S. Government has an important role in identifying what critical national infrastructure must be protected. We know that all American frontier AI labs are targets for infiltration by state and non-state actors. As the models become more capable, it is critical that frontier labs work with the U.S. Government to implement defensive measures against threat actors who would seek to abuse their models. This is why we believe there should be more robust channels between American frontier AI laboratories and the U.S. Government to facilitate threat intelligence sharing, similar to information-sharing processes used in critical infrastructure sectors, so we may shore up our collective defenses against malicious actors. Galvanizing the U.S. Government and industry capacity to sprint to prepare AI infrastructure for a world of cybersecurity AI agents is critical at this juncture.

Making Models Useful for Cyber Defenders

We therefore think a large part of making the future secure depends on our ability to make models useful for defenders and get the models into those defenders' hands. To that end, Anthropic has piloted and deployed our models with a large fraction of the world's largest cybersecurity companies, with whom we continue to partner.

We are also developing tools designed to help defenders. For example, Anthropic has released a security review tool that, with a single command, reviews a codebase for vulnerabilities and can suggest patches before code reaches production.

We envision a world where models are used by cyber defenders—in industry, Government, and by individual researchers and engineers—to secure all parts of the infrastructure that the world relies on. I am particularly encouraged by a new generation of advanced start-ups that are among the fastest and best at deploying models in creative ways to outpace attackers. We believe it is very possible that the force of innovation, spearheaded by inventive white hat companies, will be the most important factor in our ability to triumph over threat actors.

THE STAKES OF MAINTAINING U.S. LEADERSHIP IN AI

This campaign also underscores a broader strategic reality: the United States and like-minded democracies must maintain leadership in frontier AI development. Based on the current trajectory of AI development, our ability to lead at the AI frontier in the 2026–2027 time period will likely also translate directly into significant capability advancements in cyber, military, intelligence, and other critical national and economic security functions.

In this case, CCP-sponsored operators misused an American model running on American infrastructure because our technology represents the state-of-the-art. That's not a coincidence—it's a direct result of U.S. policy choices that have constrained the CCP's access to the advanced compute needed to train frontier models. Because CCP-sponsored operators had to use our systems, we were able to detect and disrupt them, and share information about the threat with the U.S. Government. That is an enormous strategic advantage.

The Trump administration has already taken important steps to advance U.S. AI leadership, including accelerating the domestic buildout of AI infrastructure, promoting Federal adoption, and strengthening safety testing and security coordination. But preserving the United States' lead in frontier AI development during this critical window depends on protecting our current advantage in compute—or the AI chips that power advanced AI systems. Restrictions on exports of advanced semi-

⁶“The Need for Transparency in Frontier AI.” Anthropic, July 7, 2025, <https://www.anthropic.com/news/the-need-for-transparency-in-frontier-ai>.

conductors and semiconductor manufacturing equipment to the CCP, building on actions initiated during the first Trump administration and expanded under the Biden administration, have been vital to preserving that edge.

Relaxing controls on advanced AI chips at this juncture could allow the CCP to close the gap in frontier AI development—producing models that may match or exceed current U.S. capabilities for cyber-offensive tasks, but without our safeguards, and using them to target U.S. critical infrastructure and national champions. Export controls on advanced semiconductors have proven effective at constraining the CCP's AI development. Without them, what any individual American company does to secure its own models becomes far less consequential. We simply won't see the attacks coming.

CONCLUSION

We are in a race against threat actors to secure systems faster and more robustly than they can be attacked. Threat actors will stop at nothing to develop, steal, or manipulate AI models to conduct increasingly sophisticated cyber attacks at scale, and we must respond urgently.

Thank you for the opportunity to appear before the committee today, and I look forward to answering your questions.

Mr. OGLES. Thank you, Dr. Graham.

I now recognize Mr. Hansen for 5 minutes to summarize his opening statement.

STATEMENT OF ROYAL HANSEN, VICE PRESIDENT, PRIVACY, SAFETY, AND SECURITY ENGINEERING, GOOGLE LLC

Mr. HANSEN. Chairmen Garbarino, Ogles, Brecheen, Ranking Members Thompson, Swalwell, Thanedar, and Members of the committee and subcommittees, thank you for the opportunity to speak with you today. My name is Royal Hansen and I serve as the vice president of privacy, safety, security engineering at Google, and, as discussed, we build the financial technology that keeps billions of people safe on-line.

As this committee knows, we stand at a critical technological inflection point. Rapid advances in AI are unlocking new possibilities for the way we work and accelerating innovation in science, technology, and beyond. Some of these same AI capabilities, however, can also be deployed by attackers, leading to understandable anxieties about the potential for AI to be misused for malicious purposes.

Until recently, our analysis showed that government-backed threat actors were using generative AI primarily for common tasks like troubleshooting, research, and content generation. Over the past year, Google's Threat Intelligence Team has identified an important shift, with adversaries not only leveraging AI for productivity gains, but deploying novel AI-enabled malware in active operations. We have identified malware families that use LLMs to generate malicious scripts, obfuscate their own code to evade detection, and use AI models to create malicious functions on demand rather than hard-coding them into the malware. This marks a new operational phase of AI abuse involving tools that dynamically alter behavior mid-execution. While still nascent, this development represents a significant step toward more autonomous and adaptive malware.

We believe not only that these highly-sophisticated threats can be countered, but that AI can supercharge our cyber defenses and enhance our collective security. LLMs can unlock new and prom-

ising opportunities, from sifting through complex telemetry to secure coding, vulnerability discovery, and streamlining operations.

Google's AI-based efforts, like Big Sleep and OSS-Fuzz, have demonstrated AI's capability to find new zero-day vulnerabilities in well-tested, widely-used software. Recently we developed CodeMender, an AI-powered agent that utilizes the advanced reasoning capabilities of our Gemini models to automatically fix critical code vulnerabilities. CodeMender scales security, accelerating time to patch across the open-source landscape. It represents a major leap in proactive AI-powered defense and includes features such as root cause analysis and self-validating patching.

We believe the private sector, governments, educational institutions, and other stakeholders must work together to maximize AI's benefits while also reducing the risks of abuse. As innovation moves forward, the industry more broadly needs security standards for building and deploying AI responsibly. That's why Google introduced the Secure AI Framework or SAIF, a conceptual framework to secure AI systems. Our recent expansion to SAIF 2.0 addresses the rapidly-emerging risks posed by autonomous AI agents and extends our proven framework with new guidance on agent security risks and controls to mitigate them.

We published a comprehensive toolkit for developers that includes resources and guidance for designing, building, and evaluating AI models responsibly. We've also shared best practices for implementing safeguards, evaluating model safety, and red teaming to test and secure AI systems. We are committed to developing technology responsibly and in a manner that is built for safety, enables accountability, and upholds high standards of scientific excellence.

For example, as part of our industry-leading security architecture, we do not offer our core products such as Search, Gmail, Maps, and YouTube in mainland China. We also do not conduct AI research, offer domestic cloud services, or have data centers in mainland China. Our comprehensive approach means we secure all components of the AI ecosystem, including data, infrastructure, applications and models.

As governments and civil society leaders look to counter the growing threat from cyber criminals and state-backed attackers, we're committed to leading the way in using AI to tip the balance of cybersecurity in favor of defenders.

Finally, this is more than a job for me. My youngest son, now 15, has suffered from a chronic illness for the past 5 years, during which time he has rarely moved from lying down in a dark, cold room. One of the few things that gives him hope is that technologies like AI and Quantum will continue to yield scientific and medical breakthroughs that will alleviate his suffering and the suffering of millions like him. Security and safety are among the critical foundations that will enable this science at digital speed. I am personally committed to that mission with the help of both the public and private sector.

We look forward to answering your questions.

[The prepared statement of Mr. Hansen follows:]

PREPARED STATEMENT OF ROYAL HANSEN

DECEMBER 17, 2025

Chairmen Garbarino, Ogles, Brecheen; Ranking Members Thompson, Swalwell, Thanedar; and Members of the Committee and Subcommittees: Thank you for the opportunity to speak with you today. My name is Royal Hansen, and I serve as vice president of privacy, safety, and security engineering at Google. Our team is responsible for building and scaling the foundational technology to keep billions of people safe on-line.

Thank you for holding this important hearing. We welcome the opportunity to provide information about Google's efforts to secure its own artificial intelligence, protect its customers' workloads, and use artificial intelligence to strengthen cyber defense and enhance our collective security.

SECURING OUR ARTIFICIAL INTELLIGENCE

Google's AI principles, published in 2018 and updated this year, describe our commitment to developing technology responsibly and in a manner that is built for safety, enables accountability and upholds high standards of scientific excellence. We have built on this work through our Secure AI Framework, as well as with extensive model hardening and various governance measures. This comprehensive approach means we secure all components of the AI ecosystem including data, infrastructure, applications, and models.

The Secure AI Framework (SAIF)

SAIF is our framework for integrating security and privacy measures into machine learning and generative AI applications and it governs how we embed controls throughout the AI system stack from data, infrastructure, application, and models. The framework, which is designed to ensure that AI models are secure by design, has six core elements:

- *Expand strong security foundations to the AI ecosystem.*—Leverage secure-by-default infrastructure protections and expertise built over the last two decades to protect AI systems, applications and users. At the same time, develop organizational expertise to keep pace with advances in AI and start to scale and adapt infrastructure protections in the context of AI and evolving threat models. For example, injection techniques like SQL injection have existed for some time, and organizations can adapt mitigations, such as input sanitization and limiting, to help better defend against prompt injection-style attacks.
- *Extend detection and response to bring AI into an organization's threat universe.*—Detect and respond to evolving AI-related cyber incidents by extending threat intelligence and other capabilities. For organizations, this includes monitoring inputs and outputs of AI systems to detect misuses, and using threat intelligence to anticipate attacks. This effort typically requires collaboration with trust and safety, threat intelligence, and counter abuse teams.
- *Automate defenses to keep pace with existing and new threats.*—Harness the latest AI innovations to improve the scale and speed of response efforts to security incidents. Adversaries will use AI to scale their impact, so it is important to use AI and its current and emerging capabilities to stay nimble and cost effective in protecting against them. It is important to remember that the vast majority of successful attacks—whether AI-enabled or not—prey on legacy systems; AI can help defenders modernize and address issues at a scale and speed that has historically proved challenging.
- *Harmonize platform-level controls to ensure consistent security across the organization.*—Align control frameworks to support AI risk mitigation and scale protections across different platforms and tools to ensure that the best protections are available to all AI applications in a scalable and cost-efficient manner. At Google, this includes extending secure-by-default protections to AI platforms like Vertex AI and Security AI Workbench, and building controls and protections into the software development life cycle. Capabilities that address general use cases, like Perspective API, can help the entire organization benefit from state-of-the-art protections.
- *Adapt controls to adjust mitigations and create faster feedback loops for AI deployment.*—Constantly test implementations through continuous learning and evolve detection and protections to address the changing threat environment. This includes techniques like reinforcement learning based on incidents and user feedback, and involves steps such as updating training data sets, fine-tuning models to respond strategically to attack attempts, and allowing the software that is used to build models to embed further security in context (e.g. de-

tecting anomalous behavior). Organizations can also conduct regular Red Team exercises to improve safety assurance for AI-powered products and capabilities. These are exactly the techniques we have used to defend Gmail, the Play Store and Chrome with AI at scale for many years.

- *Contextualize AI system risks in surrounding business processes.*—Conduct end-to-end risk assessments related to how organizations will deploy AI. This includes an assessment of the end-to-end business risk, such as data lineage, validation and operational behavior monitoring for certain types of applications. In addition, organizations should construct automated checks to validate AI performance. Nearly all businesses are increasingly digital—AI will only accelerate that trend. The controls required to mitigate risks in these processes must keep pace—some of which will be digital and some will be procedural.

Model Hardening

Our AI models are fine-tuned on large datasets of realistic attack scenarios to build intrinsic resilience. They are taught to recognize and ignore malicious instructions while still following user requests. This is, and will continue to be, an evolving space requiring rapid iterations as attackers innovate.

Over the past decade, we have evolved our approach to translate the concept of red teaming to the latest innovations in technology, including AI. The AI Red Team is closely aligned with traditional red teams, but also has the necessary AI subject-matter expertise to carry out complex technical attacks on AI systems. A core part of our security strategy is automated red teaming, where our internal Gemini team constantly attacks Gemini in realistic ways to uncover potential security weaknesses in the model. We fine-tuned Gemini on a large dataset of realistic scenarios, where automated red teaming generates effective indirect prompt injections targeting sensitive information.

Protecting AI models against attacks like indirect prompt injections requires “defense-in-depth”—using multiple layers of protection, including model hardening, input and output checks (like classifiers), and system-level guardrails. Securing advanced AI systems against specific, evolving threats like indirect prompt injection is an on-going process. It demands pursuing continuous and adaptive evaluation, improving existing defenses and exploring new ones, and building inherent resilience into the models themselves.

Securing AI Workloads

Recent headlines have highlighted several key vulnerabilities and attack vectors targeting private and public-sector entities. It is clear that legacy systems, misconfigured cloud environments, and the exploitation of known vulnerabilities remain significant concerns. Email phishing, supply chain attacks, criminal hacking, and state-sponsored cyber espionage further compound these challenges. Our approach to protecting public and private-sector entities is built on several core tenets:

- *AI-Powered Security.*—We leverage the power of AI and machine learning to enhance threat detection, automate security operations, and secure AI development.
- *Secure by Design.*—We engineer security into every layer of our infrastructure and services, from custom-designed hardware to advanced encryption techniques. To do this well requires security engineering which goes well beyond checklists and compliance requirements.
- *Zero Trust.*—We ensure that no user or device is inherently trusted, regardless of their location or network. Access is continuously authenticated and authorized based on identity, device health, and context. We developed this approach in the wake of Chinese threat actor attacks on Google over 15 years ago, and it remains as important today.
- *Shared Fate.*—We operate under a clear shared responsibility model, securing the underlying cloud infrastructure while providing tools and guidance for customers to manage their own security. We believe in a “shared fate” where our success is tied to the customer’s. We are deeply invested in the collective security outcomes of consumers, companies, and countries. We align our goals with the security and resilience of critical operations, particularly where national security is at stake.

Artificial Intelligence and Cybersecurity: Identifying Opportunities and Mitigating Risks

We stand at a critical technological inflection point. Rapid advances in AI are unlocking new possibilities for the way we work and accelerating innovation in science, technology, and beyond. Some of these same AI capabilities, however, can also be deployed by attackers, leading to understandable anxieties about the potential for AI to be misused for malicious purposes. Until recently, our analysis of gov-

ernment-backed threat actor use of AI revealed that threat actors were using generative AI primarily for common tasks like troubleshooting, research, and content generation. Over the past year, Google Threat Intelligence Group has identified an important shift, with adversaries not only leveraging AI for productivity gains, but experimenting with novel AI-enabled malware in active operations.

We have identified malware families that use LLMs to generate malicious scripts, obfuscate their own code to evade detection, and use AI models to create malicious functions on demand, rather than hard-coding them into the malware. This marks a new operational phase of AI abuse, involving tools that dynamically alter behavior mid-execution. While still nascent, this development represents a significant step toward more autonomous and adaptive malware. We have and will continue to publish on these topics, take action and enhance our products to ensure industries and societies as a whole can keep pace with the latest threats.

Today, and for decades, the main challenge in cybersecurity has been that attackers need just one successful, novel threat to break through the best defenses. Defenders, meanwhile, need to deploy the best defenses at all times, across increasingly complex digital terrain—and there is no margin for error. As we have seen in recent years, this is particularly true for legacy technology. This is the “Defender’s Dilemma,” and there has never been a reliable way to tip that balance.

Our experience deploying AI at scale informs our belief that AI can reverse this dynamic in several ways and enhance our collective security.

- AI allows security professionals and defenders to scale and accelerate their work in threat detection, malware analysis, vulnerability detection, vulnerability fixing, and incident response.
- Google’s AI-based efforts like BigSleep have demonstrated AI’s ability to find new zero-day vulnerabilities in well-tested, widely-used software. Developed by Google DeepMind and Google Project Zero, Big Sleep can help security researchers find zero-day (previously unknown) software security vulnerabilities. Since it was introduced last year, it has continued to discover multiple flaws in widely-used software, exceeding our expectations and accelerating AI-powered vulnerability research. With Big Sleep, we have demonstrated how we can find vulnerabilities that defenders don’t yet know about. In this case, we found a vulnerability that the attackers knew about and had every intention of using. We were able to detect and report it for patching before they could exploit it.
- Finding vulnerabilities is only half of the battle. Recently, we developed CodeMender, an AI-powered agent that utilizes the advanced reasoning capabilities of our Gemini models to automatically fix critical code vulnerabilities. CodeMender scales security, accelerating time-to-patch across the open-source landscape. It represents a major leap in proactive AI-powered defense and includes features such as root cause analysis and self-validated patching. This capability in particular will be the most significant security advancement in many years.

COLLABORATION TOWARD RESPONSIBLE ARTIFICIAL INTELLIGENCE ADOPTION

We believe the private sector, governments, educational institutions, and other stakeholders must work together to maximize AI’s benefits while also reducing the risks of abuse. As innovation moves forward, the industry more broadly needs security standards for building and deploying AI responsibly. That’s why Google introduced SAIF, as noted above, as a conceptual framework to secure AI systems. Our recent expansion to SAIF 2.0 addresses the rapidly-emerging risks posed by autonomous AI agents and extends our proven framework with new guidance on agent security risks and controls to mitigate them.

In addition, Google co-founded the Coalition for Secure AI (CoSAI), an open-source initiative to help all developers and deployers of AI create and maintain secure by design AI systems and help advance the framework. CoSAI helps foster a collaborative ecosystem to share open-source methodologies, standardized frameworks, and tools. Since its launch, CoSAI has made significant strides in strengthening AI security in collaboration with industry and academia in areas including Software Supply Chain Security for AI Systems; Preparing Defenders for a Changing Security Landscape; AI Security Risk Governance; and Secure Design Patterns for Agentic Systems. We have also supported the MLCommons Association’s efforts to develop AI safety benchmarks by contributing funding for the development of a testing platform, as well as technical expertise and resources. ML Commons’ shared research infrastructure helps the scientific research community derive new insights for breakthroughs in AI.

Across Google Cloud, we model and promote the adoption of responsible AI data practices that preserve our customers’ privacy and support their compliance journey.

Robust privacy commitments outline how we protect user data and prioritize privacy and the greater adoption of artificial intelligence rearms their importance. We adhere to a holistic approach to AI risk management and compliance, including focusing on employing an AI risk assessment methodology for identifying, assessing, and mitigating risks; developing and using an automated, scalable, and evidence-based approach for auditing generative AI workloads; and emphasizing human oversight and collaboration in our risk assessments and governance councils.

We use explainability tools to help understand and interpret AI predictions and evaluate potential bias; privacy-preserving technologies such as masking and tokenization and adhering to privacy laws; continuous monitoring and auditing for security vulnerabilities that AI might miss; investing in training programs to bridge the AI knowledge gap; and encouraging “interdisciplinary collaboration” between data scientists, risk analysts, and domain experts is also key.

Cybersecurity has never been a field where perfection is possible. It will remain a dynamic space for years to come, and speed and resilience will be required to defeat and contain innovative attackers. As governments and civil society leaders look to counter evolving threats from cyber criminals and state-backed attackers, we are committed to leading the way in using AI to tip the balance of cybersecurity in favor of defenders.

We appreciate the committee convening this important hearing. And we look forward to answering your questions.

Mr. OGLES. Thank you, Mr. Hansen. Just kind-of a point. First of all, thank you for sharing and I look forward to hearing more about what you’re working on, sir.

We do have votes, so we will take a short recess. I would ask all Members of the committee after the second vote to come back here as promptly as possible so that we can get to the remaining two witnesses and their opening testimony. I plan on starting as quickly as we can, if that is possible.

So thank you all. We will take a short recess.

[Recess.]

Mr. OGLES. I call to order the Committee on Homeland Security, Subcommittee on Cybersecurity Infrastructure Protection and Subcommittee on Oversight, Investigations, and Accountability will come to order.

Again, thank you, Mr. Hansen.

Then would like to recognize Mr. Zervigon for 5 minutes to summarize his opening statement. Again to the witnesses, we appreciate your patience.

STATEMENT OF EDDY ZERVIGON, CHIEF EXECUTIVE OFFICER, QUANTUM XCHANGE

Mr. ZERVIGON. Thank you. Good morning. Chairman Garbarino, Ranking Members Thompson, Thanedar, Chairman Ogles, Chairman Brecheen, and Members of the committee, thank you very much for the opportunity to testify today.

My name is Eddy Zervigon and I am the CEO of Quantum XChange. We were founded in 2018, 2 years after NIST was tasked with evaluating the algorithms to take us into the quantum age. Quantum XChange is a cybersecurity company that interoperates with the major network infrastructure vendors to enable encryption that protects data today and into the post-quantum future, with hardware and software solutions developed entirely in the United States.

While quantum computing and AI promise new breakthrough capabilities, they also introduce significant risk to our national and economic security. They must be urgently addressed. AI can enable faster, more dangerous cyber attacks, and quantum computers can

break current encryption standards, exposing sensitive data. These capabilities will be weaponized by our adversaries, creating a very dangerous imbalance in our cyber defenses.

For more than 50 years, encryption has safeguarded our data from theft and misuse. We've had the luxury of a set-it-and-forget-it mindset, trusting its strength by default. That era is now ending with quantum computing. Think about it like this. Imagine all digital communication from Government agencies sent over the past 10 years being readable by our adversaries. This is a real threat to the United States today. Rogue nation-states and state-sponsored terrorist groups are collecting encrypted data now to decrypt later with a quantum computer.

Further, now imagine our adversaries reading sensitive Government data in real time and altering it without anyone knowing. This could be tomorrow's reality. Public and private-sector work on quantum resilient solutions is on-going. Technologies like post-quantum cryptography, PQC, or quantum-safe encryption algorithms are part of the solution, but not the complete answer.

Despite our best efforts, post-quantum cryptography may still be vulnerable to quantum-related attacks. All of which raises the fundamental question and challenge what happens when an algorithm breaks? Because it is a when and not if. Every agency CIO, enterprise CISO, security vendor, and network gear manufacturer must be able to answer that question.

In our view, what's needed to ensure data security and confidentiality in the quantum age is an architectural approach, not just a new algorithm. This architectural approach enables agencies to focus on securing the network that data travels on to strengthen the existing infrastructure against quantum attacks while minimizing disruption to existing operations. This is how our Government agencies need to be protected.

When you have valuables in your house, the first step isn't going out and buying a new jewelry box with biometric access controls. It's locking your front and back doors so the house is secure and harder to get in. Once your home is secure, then you can figure out what specific rooms need further locks or security measures to protect your valuables and sensitive documents.

Federal agencies handling sensitive data need to act now and follow the lead set by Customs and Border Protection. Our work with CBP to incorporate PQCs across their network infrastructure in 2026 has shown that you can begin to secure your networks today with quantum-resistant technologies in a FIPS-validated way without having to rip and replace your entire infrastructure.

I cannot stress enough the timing here is critical. Agencies that fail to prepare today risk leaving their data vulnerable. Every day that we are not quantum-resistant is another day that data is harvested to be decrypted later.

It is important to note that we at Quantum XChange are not the only ones advocating for action today. The Quantum Industry Coalition of which we are part of, as well as Amazon Web Services, Google, IBM, Microsoft, Accenture, and others, believe that agencies handling sensitive Government data should be actively working and preparing for the transition and should begin migrating to high-risk systems to FIPS/NIST validated PQC where possible.

Having the opportunity to meet with several of your offices, I was often asked what can Congress do? Through this committee's leadership and building off the work previously done, Congress can accelerate the time lines for PQC compliance, allocate the budget to allow migration process to begin, and work with leaders within the administration to encourage adoption as the technology is readily available and deployable today.

America's defenses cannot stop at our physical borders. Through your leadership and efforts and in partnership of private-sector partners, like us, we can and secure—we can and will secure America's digital borders, too.

In closing, I want to thank you again for the opportunity to offer some thoughts today and I look forward to your questions. Thank you.

[The prepared statement of Mr. Zervigon follows:]

PREPARED STATEMENT OF EDDY ZERVIGON

DECEMBER 17, 2025

Good morning, Chairman Garbarino, Ranking Member Thompson, Chairman Ogles, Chairman Brecheen, and Members of the committee. Thank you very much for the opportunity to testify today.

My name is Eddy Zervigon, and I am the CEO of Quantum XChange. We were founded in 2018, 2 years after NIST was tasked with evaluating the algorithms to take us into the quantum age. Quantum XChange is a cybersecurity company that interoperates with the major network infrastructure vendors to enable the encryption that protects data today and into the post-quantum future with hardware and software solutions developed entirely in the United States.

While quantum computing and AI promise new breakthrough capabilities, they also introduce significant risks to our national and economic security that must be urgently addressed. AI can enable faster, more dangerous cyber attacks and quantum computers can break current encryption standards, exposing sensitive data. These capabilities will be weaponized by our adversaries, creating a very dangerous imbalance in our cyber defenses.

For more than 50 years, encryption has safeguarded our data from theft and misuse. We've had the luxury of a "set it and forget it" mindset, trusting its strength by default. That era is now ending with quantum computing.

Think about it like this: Imagine all digital communications from Government agencies sent over the past 10 years being readable by our adversaries. This is a real threat to the United States today; rogue nation-states and state-sponsored terrorist groups are collecting encrypted data NOW to decrypt later with a quantum computer.

Further, now imagine our adversaries reading sensitive Government data in real time, and altering it without anyone knowing. This could be tomorrow's reality.

Public and private-sector work on quantum-resilient solutions is on-going. Technologies, like post-quantum cryptography (PQC) or quantum-safe encryption algorithms, are part of the solution but not the complete answer. Despite our best efforts, post-quantum cryptography may still be vulnerable to quantum-enabled attacks.

All of which raises this fundamental question and challenge: What happens when an algorithm breaks (because it is a when, not if)? Every agency CIO, enterprise CISO, security vendor, and network gear manufacturer must be able to answer that question.

In our view, what's needed to ensure data security and confidentiality in the quantum age is an architectural approach, not just a new algorithm.

This architectural approach enables agencies to focus on securing the network that data travels on to strengthen the existing infrastructure against quantum attacks, while minimizing disruption to existing operations. This is how our Government agencies need to be protected. When you have valuables in your house, the first step isn't buying a new jewelry box with biometric access controls, it's locking your front and back doors, so the house is secure and harder to get in. Once your home is secure, then you can figure out what specific rooms need further locks or security measures to protect your valuables and sensitive documents.

Federal agencies handling sensitive data need to act now and follow the lead set by Customs and Border Protection. Our work with CBP to incorporate PQCs across their network infrastructure in 2026 has shown that you can begin to secure your networks today with quantum-resistant technologies in a FIPS-validated way, without having to rip and replace your entire infrastructure. I cannot stress enough that timing here is critical.

Agencies that fail to prepare today risk leaving their data vulnerable. Every day that we are not quantum-resistant is another day that data is harvested, to be decrypted later. It is important to note, that we at Quantum XChange are not the only ones advocating for action today. The Quantum Industry Coalition, which we are a part of and includes Amazon Web Services, Google, IBM, Microsoft, Accenture, and others believes “that agencies handling sensitive government data should already be actively preparing for the transition and should begin migrating high-risk systems to FIPS/NIST validated PQC where possible.”

Having the opportunity to meet with several of your offices, I was often asked “What can Congress do?” Through this committee’s leadership, and building off the work previously done, Congress can accelerate the time lines for PQC compliance, allocate the budget to allow the migration process to begin, and work with leaders within the administration to encourage adoption, as the technology is readily available and deployable today. America’s defenses cannot stop at our physical borders. Through your leadership and efforts, and in partnership with private-sector partners like us, we can and will secure America’s digital borders too.

In closing, I want to thank you all again for the opportunity to offer some thoughts today and look forward to your questions.

APPENDIX.—QUANTUM INDUSTRY COALITION POSITION ON POST-QUANTUM CRYPTOGRAPHY

OCTOBER 23, 2025

The National Institute of Standards and Technology (NIST) has approved the first set of postquantum cryptographic (PQC) algorithms, in what promises to be an iterative process moving forward. NIST has been leading the migration charge for close to a decade, evaluating and approving the algorithms and delivery architectures that will protect our data networks into the post-quantum era.

The Federal Government has set time lines for the adoption of these post-quantum algorithms through Legislation and Executive Orders. Government agencies should already be preparing for PQC transition through education, cryptographic inventory, risk assessments, transition strategies, and pilots. At the same time, the ecosystem of innovative start-ups and established players surrounding the delivery of these algorithms has progressed to a point where transition is possible in some high-risk areas, such as securing the network layer.

It is our position that agencies handling sensitive Government data should already be actively preparing for the transition and should begin migrating high-risk systems to FIPS/NIST validated PQC where possible.

Quantum Industry Coalition Members Include:

Accenture
D-Wave
Entanglement Institute
IonQ
Quantinuum
Rigetti Computing
Xanadu
Amazon Web Services
Cold Quanta
Dirac
Google
MesaQuantum
Quantum Corridor
SandboxAQ
Anametric
EeroQ
IBM
Microsoft
Quantum Machines
SEEQC
Atom Computing
enQase

Inflection
Qolab
Quantum XChange
Strangeworks

Mr. OGLES. Thank you, Mr. Zervigon.

I now recognize Mr. Coates for 5 minutes to summarize his opening statement.

**STATEMENT OF MICHAEL COATES, FOUNDING PARTNER,
SEVEN HILL VENTURES**

Mr. COATES. Chairman Ogles, Ranking Member Swalwell, Chairman Brecheen, and Ranking Member Thanedar, thank you for the opportunity to testify. I'm honored to be here to discuss the changing cybersecurity landscape and the impacts of artificial intelligence and quantum computing. My perspective is grounded in over 20 years of experience in cybersecurity, including service as a chief information security officer, leadership in global software security organizations, founding a technology start-up, and investing in cybersecurity innovation.

Today we sit at the precipice of significant change. While much attention is paid to AI and future breakthroughs like AGI, the most immediate impact on cybersecurity is not the creation of entirely new threats. Instead, AI and quantum technologies are collapsing the time, cost, and skill required to conduct cyber operations. These changes are outpacing existing technical, regulatory, and operational defenses, fundamentally reshaping the threat landscape.

Historically, different attackers, nation-states, cyber-criminal organizations, and lone hackers were constrained by skill, resources, and scale. The most sophisticated attacks were largely limited to nation-states, while criminals focused on repeatable, monetizable techniques. That constraint is rapidly changing. Recent real-world examples, such as the report issued by Anthropic, show AI systems being used as a central orchestration layer for complete cyber operations, coordinating reconnaissance, exploitation, and execution with limited human involvement. While the techniques themselves may not be novel, the orchestration and automation represent a meaningful shift in adversary capability.

Agentic AI further removes human constraints. Autonomous systems are not limited by time, fatigue, or tension, and research recently released from Stanford, Carnegie Mellon, and Grace One AI already show AI-driven penetration testing performing at or above the level of highly-skilled professionals at a fraction of the cost.

At the same time, AI is accelerating vulnerability discovery and exploitation. AI-powered software analysis is capable of identifying previously-unknown zero-day vulnerabilities faster than ever. Yet for many organizations, the long-standing challenge has not been awareness that a vulnerability exists, but rather the inability to patch and remediate quickly. As attack time lines compress, this operational inertia becomes more dangerous.

The practical result is a dramatic reduction in the time available for defenders. Comprehensive attacks are easier to launch, the pool of capable adversaries expands, and smaller organizations, such as hospitals, schools, and small businesses are increasingly exposed to

the same level of adversarial capability once reserved for critical national infrastructure.

This compression of time changes the nature of cyber risk itself. Defenders are often no longer responding to early indicators, but to attacks that are already in progress. Intelligent automation allows attacks to become continuous rather than episodic, eroding assumptions that organizations can recover between incidents or rely on periodic assessments.

The widening gap between machine-speed attacks and human-speed defenses means cybersecurity outcomes are increasingly determined by whether defenses can operate at comparable speeds. These shifts have clear implications for defense policy and coordination.

First, secure-by-design principles must become a baseline expectation, particularly as AI increasingly writes and modifies software. Second, regulatory clarity is critical. Fragmented or ambitious regulations can slow defensive responses in an environment or speed matters. Third, public-private coordination remains essential, ensuring that defensive learning keeps pace with adversarial innovation. Fourth, defensive capabilities must increasingly rely on automation and autonomy as purely human-driven defenses will struggle to keep up. Fifth, finally, quantum preparedness is necessary. While post-quantum cryptographic standards exist, the challenge lies in the time and coordination required to migrate existing systems before an adversary achieves cryptographically-relevant quantum capability.

Finally, trust and transparency in AI systems are crucial. AI reflects the data, incentives, and governance under which it is trained. In a security-related context, understanding potential model bias and model origin is as important as performance.

Artificial intelligence and quantum computing are accelerating forces that dramatically reshape cybersecurity. Our success will depend on whether our technical, operational, institutional responses can adapt at a comparable pace.

Thank you and I look forward to your questions.
[The prepared statement of Mr. Coates follows:]

PREPARED STATEMENT OF MICHAEL COATES

DECEMBER 17, 2025

Chairman Ogles, Ranking Member Swalwell, Chairman Brecheen, and Ranking Member Thanedar, I thank you for the opportunity to testify before you today. I'm honored to be here to speak about the changing landscape in cybersecurity and the resulting impacts from AI and quantum computing.

The perspective I will share is grounded in over 20 years of experience in cybersecurity, including service as a chief information security officer, a chairman of a global non-profit advancing the state of application and coding security, a technology start-up founder, and a venture capital investor supporting cybersecurity innovation.

Today we sit at the precipice of significant change. While advancements in AI and development toward AGI are widely discussed, the practical and operational impacts to cybersecurity defenders are less often examined.

The fundamental reality is not that AI and quantum are creating new types of threats, but rather they are collapsing the time, cost, and skill required to conduct cyber operations. These changes are outpacing the existing technical, regulatory, and operational defenses. This shift reshapes the cyber threat landscape and forces a reconsideration of how we defend critical systems in an era defined by speed, automation, and intelligent scale.

WHAT IS CHANGING: THE COMPRESSION OF CYBER CAPABILITY

CAPABILITY COMPRESSION & ORCHESTRATION EXPANDS THE ATTACKER BASE

Corporations and citizens potentially face a variety of threat agents including highly-funded nation-state adversaries, financially-motivated cyber-criminal organizations, and lone hackers motivated by ideology. Each attacker type has different skills and resources at their disposal and to date, these have constrained the complexity or scale of cyber attacks available to each adversary.

The most advanced attacks were often only launched by nation-state adversaries against select targets. Whereas cyber-criminal entities focused their efforts on pipelines of optimized offensive security services, such as ransomware extortion, to monetize the compromise of businesses or individuals.

Robust security attacks require a series of steps spanning reconnaissance, exploitation, command and control, and delivery of the ultimate objective, such as data theft or system modification. Each of these components could be executed by a well-funded nation-state adversary or a competent cyber-criminal organization, but it was not as achievable for the lone hacker or unsophisticated security hacker. This is rapidly changing.

As demonstrated in the November, 2025 Anthropic report “Disrupting the first reported AI-orchestrated cyber espionage campaign”,¹ a nation-state adversary used AI systems as a central brain and point of coordination for a complete security attack against multiple targets across the United States. AI was used to execute and interpret results for each step of the attack and as an overall orchestration layer, with the human adversary only interacting at a few decision points.

While this attack may not have demonstrated new or novel attack methods, the orchestration and use of AI is a critical development in the ecosystem of the cybersecurity adversary.

AGENTIC ATTACKS REMOVE HUMAN CONSTRAINTS

Agentic AI systems will enable the attacker to no longer be bound by time of day, hours awake, or the need for food or sleep. Autonomous agentic systems are replicating the most advanced attackers and will be able to target with accuracy and ease.

This is no longer theoretical as research just released by Stanford² shows that an autonomous AI penetration-testing agent already performs at or above the level of most highly-skilled professional security testers, outperforming 9 out of 10 participants in a live network test with an 82 percent valid vulnerability discovery rate at a fraction of the cost.

ACCELERATION OF VULNERABILITY DISCOVERY AND EXPLOITATION

Furthermore, the increasing power of AI for software vulnerability analysis is enabling faster and more accurate detection of previously-unknown zero-day security vulnerabilities. For example, Google’s Big Sleep, a collaboration between Google Project Zero and Google DeepMind, has discovered a critical zero-day vulnerability in the major software SQLite Database Engine.³

Over the past decades, the challenge for many organizations has not been knowledge that a vulnerability existed, but rather the operational inertia to deploy, test, and productize the software patch. In fact, the 2025 Verizon Data Breach Investigations Report found that vulnerability exploitation was the initial access vector in 20 percent of breaches, and that defenders often cannot remediate fast enough—organizations fully remediated only about 54 percent of vulnerabilities in network edge devices, with a median remediation time of 32 days, while CISA KEV vulnerabilities can be mass exploited in a median of 5 days.⁴

THE PRACTICAL RESULT: REDUCED TIME FOR DEFENDERS

With AI orchestration, the ease of launching comprehensive cybersecurity attacks against any target is substantially reduced. The result is that many more potential adversaries now have the means to execute these attacks.

¹<https://www.anthropic.com/news/disrupting-AI-espionage>.

²<https://arxiv.org/pdf/2512.09882>.

³<https://cloud.google.com/blog/products/identity-security/cloud-ciso-perspectives-our-big-sleep-agent-makes-big-leap>.

⁴<https://www.verizon.com/business/resources/reports/dbir/>.

In addition to an increase in attacks against the most critical targets, this development will also result in lesser-profile targets, such as small businesses across the country, being subjected to full-scale security assaults.

The direct result of this change will be a dramatic drop in the time available for defenders to detect attacks, initial compromise, or lateral movement before critical access or sensitive data is breached. Taken together, these shifts do not just increase cyber risk, they fundamentally change the speed at which cyber incidents unfold.

WHY TIME COMPRESSION CHANGES THE NATURE OF CYBER RISK

The compression of time, cost, and skill required to conduct cyber operations fundamentally changes how cyber risk manifests in practice. While individual techniques may appear familiar, the speed at which attacks now unfold alters the balance between attackers and defenders in ways that existing security models were not designed to accommodate.

The most immediate consequence is a dramatic reduction in the time available for defenders to detect and respond to malicious activity. AI-enabled orchestration and automation allow attackers to move from initial access to lateral movement and impact far more quickly than in the past. In many cases, defenders are no longer responding to early indicators of compromise, but to attacks that are already well under way.

This compression of time disproportionately affects organizations that lack large, specialized security teams. While highly-resourced enterprises may be able to invest in advanced detection and response capabilities, smaller organizations, including hospitals, schools, food processing facilities, and small businesses often rely on delayed or manual processes. As sophisticated attacks become easier to launch and less expensive to operate, these lower-profile targets increasingly face the same level of adversarial capability once reserved for critical national infrastructure.

At the same time, intelligent automation and scaling by adversaries is shifting the risk of attacks from periodic events to a continuous threat. AI-driven attacks do not require sustained human attention and can operate persistently, adapting to defenses and retrying failed approaches automatically. This erodes traditional assumptions that organizations can recover between incidents or rely on periodic assessments to maintain security.

Existing defensive and governance models further compound this challenge. Over the past decades, many major breaches did not occur because vulnerabilities were unknown, but because organizations were unable to deploy patches or mitigations quickly enough. As AI accelerates vulnerability discovery and exploitation, this operational inertia becomes more consequential. The gap between awareness and action grows more dangerous as attack time lines compress.

The result is a widening gap between the speed and accessibility of modern cyber attacks and the ability of most organizations to respond. As AI compresses attack time lines and expands the pool of capable adversaries, cybersecurity outcomes will increasingly be determined by whether defenses can operate at machine speed.

IMPLICATIONS FOR CYBER DEFENSE, POLICY, AND COORDINATION

The advancements in artificial intelligence and quantum computing present significant opportunities for innovation, but without appropriate alignment between technology, operations, and governance, they also introduce material cybersecurity risk. The shifts described earlier are not theoretical, and they cannot be addressed by any single organization or sector acting alone.

The following are key areas where attention is warranted to increase the cybersecurity posture of our organizations and critical systems.

Secure by Design as a Baseline Expectation

As software is increasingly written, analyzed, and modified by AI systems, secure design principles must be integrated into the creation of software from the outset. Initiatives such as CISA's Secure by Design program, along with industry standards promoted by organizations like OWASP and the Cloud Security Alliance, provide important guidance. Supporting these organizations and reinforcing these efforts helps ensure that speed and automation do not come at the expense of security fundamentals.

Regulatory Clarity That Supports Speed and Innovation

Clear and transparent regulatory frameworks are necessary to enable rapid innovation while maintaining responsibility for security and safety. In an environment where threats evolve quickly, ambiguity or fragmentation in regulation can unintentionally slow defensive response and increase systemic risk. Policy should seek to

provide clarity and consistency without constraining the ability of organizations to adapt at machine speed.

Public-Private Coordination on AI-Driven Cyber Threats

The pace of change in the cyber threat landscape reinforces the importance of strong public-private partnerships. Effective coordination, information sharing, and joint response mechanisms help ensure that defensive learning keeps pace with adversarial innovation. These partnerships remain a critical component of national cyber resilience as AI-driven threats continue to evolve.

Migration Toward Autonomous Defensive Capabilities

As attackers increasingly rely on automation and agentic systems, purely human-driven defenses will struggle to keep pace. Continued investment in research, development, and deployment of intelligent and autonomous defensive systems is necessary to address machine-speed threats. This includes supporting innovation across both the public and private sectors.

Quantum Preparedness for Cryptographic Systems

Stable, cryptographically-relevant quantum computing would render many of today's widely-deployed public-key encryption algorithms ineffective, impacting secure communications across government, industry, and critical infrastructure. While post-quantum cryptographic standards already exist, the primary challenge is the time and coordination required to migrate existing systems. Deliberate preparation is crucial to avoid a reality where an adversary achieves cryptographically-relevant quantum capabilities first and thus access not only to future communications, but potentially to sensitive data captured and stored today.

Trustworthiness and Transparency in AI Systems

As AI systems are increasingly embedded into security-sensitive workflows, trust in operation becomes crucial. Large language models reflect the data, incentives, and governance structures under which they are trained, and these factors can materially influence reliability and security outcomes.

Bias in AI systems—whether intentional or unintentional—can affect how software is generated, how alerts are prioritized, and how decisions are made. In security-critical contexts, performance alone is not sufficient; the provenance, training, and oversight of AI systems must also be considered as part of risk assessment.

Furthermore, greater transparency in software procurement and composition is needed. Requiring bill of materials and software contracts to disclose the use of AI within software, as well as the specific models and model origins, can help organizations better assess risk and make informed security decisions, particularly in sensitive or critical environments.

Artificial intelligence and quantum computing are accelerating dynamics that dramatically shift the cybersecurity landscape. As AI and quantum computing continue to advance and are increasingly leveraged by cyber adversaries, success will depend on whether our technical, operational, and institutional responses can adapt at comparable pace.

I appreciate the opportunity to share these observations and look forward to your questions.

Mr. OGLES. Thank you, Mr. Coates.

Members will be recognized by order of seniority for their 5 minutes of questioning. I now recognize myself for 5 minutes.

Dr. Graham, Anthropic's investigation into the recent PRC-affiliated cyber incident involving Claude suggests we may be approaching a turning point in how cyber operations are conducted, where AI systems, once asked—tasked by human operators, can execute and refine large portions of a cyber attack at machine speed rather than human speed. Obviously you touched on this in your opening statement—should this incident be understood as an early warning of the future of AI systems, how they are autonomously, you know, writing and adapting to systems? Quite frankly, from a defensive perspective, you know, what capability gaps do we have? Where do we need to be anticipating?

I mean, I see a horizon that we can't quite define because of the rapidness and just the evolving nature of the technology. I go back

to kind-of the arms race. There was a point at which, between the United States and Russia, there was this detente, there was this, you know, mutually-assured destruction, where it was at some point we all had enough nukes to kill everybody and blow the whole world up. It was all about delivery systems at that point.

AI is different. There is no horizon. There is no kind-of point at which I think it stops, that there is a ceiling. So, please, take it away.

Mr. GRAHAM. You're correct that we are at a change point. There are a couple of change points here. The first that we see now is, to our understanding, this is the first time where these models will now be sought and used by sophisticated state actors. We've been tracking this trend line for many years. This is the clearest evidence for the first time that this is now happening.

But it's also possible this gets more serious and the stakes become much higher. As you say, it's very possible that attacks from here on might scale if we don't properly secure and safeguard the models. It's also possible that while in this case we didn't see an instance of novel—or novel methods of attack, it's very possible that models could get that good.

What's important now is a few things. First, it's really hard to win if we can't see the playing field. I think the easiest way to start is continuing to evaluate the capabilities of these models. This is something industry should do, this is something Government should do. Second, we should be sharing threat intelligence as it happens so that we can mitigate as fast as possible. Third, as you say, we need to make sure defenders have the advantage, particularly the United States, make sure that it defends itself faster than it can be attacked. We are working very hard, and I think all industry needs to work hard to make that happen.

Mr. OGLES. You know, a follow-up onto that point, you know, clearly, when you look at the investments that China is making on these quantum capabilities, AI, et cetera, you know, there is a requirement between, you know, their private sector, if you want to even call it a private sector because most of it is state-owned, that any innovation is immediately shared with the State. So, as you mentioned, there is going—for us to be successful, there is going to have to be this collaboration between private and Government, quite frankly.

But one of the things, and obviously that is easier to accomplish, but there is—I foresee a need where the industry itself is going to have to be sharing information. Of course, the problem you get into there is the proprietary nature of things. You know, obviously there is the monetization factor that comes into that. But at the end of the day, we are talking about the homeland. So how do you see that working in practice, understanding the complications that we have essentially in a free market?

Then another layer to that is essentially the Five Eyes, the Seven [sic] Eyes, our European partners, who are aligned with us in our values, who understand the existential threat that China poses. Again, it is important for everyone to understand that China is probing us daily to look for weaknesses and opportunities to take advantage of information that is not properly secured. What is dif-

ferent about this is the leveraging and the scale and the percentage, if you will, that AI was leveraged. Sir.

Mr. GRAHAM. It is very, very important that industry does share the information that it has between itself. It's very important it shares that with Government. It's very important that industry develop solutions now, whether it's by improving the models or building tools and putting them in the hands of the defenders. I think just making the models good enough isn't sufficient. We need to make sure people are using it to proactively defend critical infrastructure. One way that I think Government can be extremely helpful here is identifying the critical infrastructure that needs to be defended in this new era of cybersecurity and allowing industry to point out its talents and innovation toward that.

Mr. OGLES. Well, I want to thank, again, all of you for being here and, quite frankly, to Anthropic for your report. I think it was one of those inflection points that we all understood the seriousness of this. But your report, I think, really put a light on where we are at and where—and some of our vulnerabilities.

I now recognize the Ranking Member, the gentleman from Michigan, Mr. Thanedar, for his 5 minutes of questions.

Mr. THANEDAR. Thank you, again, Chairman Ogles. Appreciate all of our witnesses.

You know, I remain deeply worried and concerned about President Trump's decision to allow export of advanced chips to China. I just don't understand other than his desire to please a donor. I just don't understand, why would we give such advanced technology to an adversary like China, who can then use this technology to attack us? Who could use this technology to cyber attack our critical infrastructure?

Dr. Graham, how would China having access to this advanced chips, how will that help advance their AI technology? Will that pose a threat to the United States, our national security?

Mr. GRAHAM. We view it as, first, extremely important that America retains its AI leadership. The most important input to this is the compute advantage. My concern from watching these models progress in their capabilities, especially as a result of the cyber espionage campaign, is that if Chinese frontier labs have access to similar amounts of compute, they could train models that are equally or more capable in the cyber domain and that this could unleash new scale and new sophistication, and we will have a harder time detecting and defending it.

Mr. THANEDAR. Thank you. Thank you. I want to shift my focus. I only have a limited time. I want to shift my focus on immigration.

You know, in his first term, President Trump's first term, and now in his second term, there is just so much of hate against immigrants. Yet we know, and I hope the panel agrees with me, that the United States technology industry has benefited greatly from immigrants.

Just by answering yes or no from the witnesses, does your companies have immigrants, skilled immigrants, and do you depend on them? Yes, no?

Mr. GRAHAM. Anthropic is composed of many of the best talent from around the world.

Mr. THANEDAR. Anybody thinks we should have less of skilled immigrants on the panel here? Should we restrict access of immigrants to our technology companies, immigrants who help us keep on the edge?

Well, certainly, you know, I am myself an immigrant. Twenty-four years old, I came here escaping poverty in India, got a Ph.D. in chemistry, became a serial entrepreneur, ran many pharmaceutical companies, developing technology that helped us stay on top of innovation. You know, while it is important that American jobs be protected, it is important that we create skills. But at the same time, our tech industry heavily depends on skills, skill sets, immigrant skill sets.

Have the actions of the Trump administration—how has acts of the Trump administration made it difficult to retain international talent in your companies with regard to both international workers choosing to leave or being forced to leave due to discrimination changes, the hardship that they have in terms of getting their status adjusted, getting their green cards, the long delay in processing, making it harder to get an H-1B visa? I just wanted to understand what kind of impact these administration's positions are doing to your ability to grow your companies, grow your new technology for the United States. Anybody? Yes.

Mr. GRAHAM. Well, it's not my issue area that I cover in the company. Speaking for my team, it's really important that I find and hire the best people around the world that are committed to our mission of making AI stay secure and ensuring America's leadership.

Mr. THANEDAR. Yes. Anybody else? How important is immigration?

Mr. HANSEN. I mean, I'd just say again, it's not—you'd have to talk to our H.R. department so we can come back to you with—you know, I'll relay that question to the teams.

Mr. THANEDAR. What percent of your organization has immigrants?

Mr. HANSEN. I wouldn't know the exact number, but certainly we do have green cards and immigrants that work at Google.

Mr. THANEDAR. Thank you. Anybody else?

You know, again, the need continues and for us, America, to have its edge on innovation, whether it is cybersecurity, AI, quantum, we must have skilled work force. If that means we have to depend on immigrants, so be it.

Thank you. I yield back.

Mr. OGLES. The gentleman yields back.

I recognize the Chairman of the Subcommittee on Oversight, Investigations, and Accountability, the gentleman from Oklahoma, Mr. Brecheen.

Mr. BRECHEEN. Thank you Mr. Chairman.

Mr. Hansen, just before I get started, prayers over your son. May the Lord do what human hands can't. Appreciate your passion, appreciate your vulnerability in sharing that.

Also appreciate what you expressed about limiting services for mainland China. I think that is great that your company is willing to do that. My hope is that others would watch your concern over

proprietary information and desire to make sure that U.S. citizenry is protected and follow your lead.

Mr. Graham, you talked about that you felt like that robust intelligence sharing could be enhanced. So what is it that you are seeing that could be improved upon about, of course, your front line, the free market, Government learns from it? What can the Fed be doing to a greater level, Homeland Security specific to this committee's assignment, to make sure that that robust intelligence sharing is happening, so that, you know, in real time we are sending out information that others can be protected based upon immediate experience?

Mr. GRAHAM. Yes. A fundamental issue here is that as the technology gets better, we're going to start seeing new patterns that are potentially more sophisticated than either in industry or across government we have not seen before in terms of what these attacks look like. I think the first most important thing is we need good and quick and sensitive channels to share the novelty of this information, possibly within and to Government and cross industry.

We probably need to get ahead of it as well. So we need to be able to share information prior to the attack occurring. We regularly brief and share information about model capabilities as they're advancing. In general, any effort here I think is extremely valuable and I think is going to put all of industry in a better position.

Mr. BRECHEEN. Yes. One of the things we can do is there are people that work behind the scenes that never, you know, get in front of the limelight of Government. So without naming names, what division with Homeland Security can we highlight to just send a special thank you to working with you?

Mr. GRAHAM. I'm not an issue expert in the specific components of homeland security, but would very happily follow up with you to talk more.

Mr. BRECHEEN. That would be great. We want to make sure we're congratulating those groups that are taking your experience seriously.

I want to talk about the at-scale capability of 80 to 90 percent of nonhuman hands on what would be formally labor-intensive, now turned into generated by computer processing. So Mr. Hansen, if AI is utilized to provoke, then AI can be utilized to defend. So how can we enhance our scale of utilizing AI to wall off?

Mr. HANSEN. It's exactly the right question. So when you talk about what we can do is I think of the old adage about the cobbler's children who don't have shoes. So there are far more defenders in the world than there are attackers. But we need to arm them with the—that same type of automation that you saw in the attack described by Anthropic. Because it's just in many ways using commodity tools that we already have to both find and fix vulnerabilities. Those can be turned from offensive capabilities to the patching and fixing. But the defenders have to put shoes on. They have to use AI in defense.

So while the attackers are experimenting, we need the defenders to be experimenting and becoming great users of AI to find the same vulnerabilities that were described, but instead of exploiting them, to patch them. That's the kind of—I mentioned CodeMender

is our project, which takes advantage of this, you know, vibe coding, if you want to call it. It's easier and easier to code. We make it easier and easier to patch.

With so much of our problems based on legacy technology, small companies, others, that's the only way we're going to get ahead. This defender's dilemma of attacker needs to be right once, defender needs to be right all the time, AI can help the defender be right all the time. That's what we need to do.

Mr. BRECHEEN. Mr. Zervigon, if I did a horrible job of pronouncing your name, you have a last name like mine, I apologize. Mr. Coates, you have taken the time to be here. I have got 30 seconds. If there is anything, because this is such an exploratory exercise for so many of us that are not experts, is there anything you want to just highlight? I have got 20 seconds to split between the two of you.

Mr. ZERVIGON. I would say innovative results demand innovative time lines. Right? You can't be operating on legacy time lines in order to achieve innovative results to protect the homeland.

Mr. COATES. The piece I would add is that the information sharing is critical. Staying abreast of how this is evolving is going to be one of the most important pieces amongst enterprises fighting against the new threats.

Mr. BRECHEEN. I look forward to highlighting Homeland Security staff with our committee staff.

Thank you, Mr. Chairman.

Mr. OGLES. The gentleman yields back.

I recognize the gentleman from Rhode Island, Mr. Magaziner.

Mr. MAGAZINER. Thank you, Chairman.

I am going to get right to the point. The Chinese government just launched the first-ever AI-powered cyber attack against our country that we know of. At the same time, President Trump is selling the powerful H200 Nvidia chips, the next generation chips, to China. I will ask any of our 4 experts, does anybody think this is a good idea? Or our colleagues or anyone, does anyone want to defend this decision?

Like they are literally—they are engaging in cyber warfare against us right now. They just did it. They just launched the first AI-powered cyber attack against U.S. organizations. Why in the world, given that they just did this, what, a couple months ago, would we be giving them these next generation chips now? At the very least, we ought to be holding them back until we have some way of verifying that these chips are not going to be used to attack us.

So I will ask again. Any of our witnesses, Mr. Graham, Mr. Coates, anyone, why is it concerning to you that China is about to receive these H200 chips from Nvidia? Mr. Coates, would you like to take a stab at it?

Mr. COATES. The defenses that we put into our LLMs, that Anthropic, that Google, and others are doing to provide safety, are things that we can control and we can use to prevent future type attacks from China using these resources. As China achieves the same capabilities and their technology from these chips, we lose control of the ability to put those safeguards in place and we're on our heels. So I agree with the concern that's being raised.

The other piece that I will mention here is that as China provides greater frontier models, like DeepSeek, and it's appealing to U.S. software corporations to integrate that into their stack for performance regions, we have to remember that that is essentially delegating decision making and trust to China, even though it might be U.S. software. We need greater focus on that.

Mr. MAGAZINER. Yes, I mean, look, cybersecurity is a bipartisan issue. I believe that there are people on both sides who care genuinely about keeping us safe in the cyber domain. But, like, I don't know how anybody can be OK with this chip sale given what literally just happened 2 months ago. That is something that I think we need to find a way as a Congress to deal with because the administration, I fear, has made a grave mistake.

I want to talk about the attack more specifically because we need to learn as much as we can from it. Mr. Graham, I am grateful that Anthropic was able to detect and then report about the nature of the attack, but my understanding is it took about 2 weeks for Anthropic to realize that the attack was happening, give or take. Is that correct? Can you explain to us, you mentioned it in your written testimony, can you explain to us generally why it took so long and what lessons you have learned, and how you can now detect similar attacks, hopefully faster in the future?

Mr. GRAHAM. Yes. The first thing to note is we ultimately did detect and disrupt the attack. When we did, it was clear that this was a highly-resourced, sophisticated effort to get around the safeguards in order to conduct the attack. Very specifically, what they did was they used a private obfuscation network to ensure that it was difficult to trace where the operations were coming from. They broke out the attack into small components that individually looked benign, but taken together form a broad pattern of misuse. Then ultimately, they deceived the model into believing that it was performing ethical—I mean—

Mr. MAGAZINER. They basically told the model, help us figure out how to protect ourselves from a cyber attack, but, in so doing, the model revealed the vulnerabilities to a cyber attack. Is that, in layman's terms, what happened?

Mr. GRAHAM. That is one of the components. That's—it's one of the key issues with cybersecurity.

Mr. MAGAZINER. Yes. I mean, I would just say as like a layperson, that that seems like something that, you know, ought to be flagged. Right? If someone says, help me figure out what my vulnerabilities are, there should be an instant flag that someone may actually be looking for vulnerabilities for a nefarious purpose.

So I will just ask for the time I have left to any of our witnesses, I mean, what regulation is required to ensure that commercially-available AI products have adequate guardrails in place? We appreciate the, you know, the efforts that companies are already undertaking, but there should be some sort of a baseline of standards that we set as a country, should there not?

Mr. HANSEN. We released this Secure AI Framework, SAIF, and then there's a 2.0 version, as well as a Coalition for Secure AI where we're not just helping set standards, but open source the implementations so, broadly, people can take advantage of and use those in their infrastructure.

Mr. MAGAZINER. All right, thank you all. I yield back.

Mr. OGLES. The gentleman yields back.

I now recognize the gentleman from Texas, Mr. Luttrell.

Mr. LUTTRELL. Thank you, Mr. Chairman.

Mr. Zervigon, did I say that right?

Mr. ZERVIGON. Perfect. Yes, sir.

Mr. LUTTRELL. You spoke on architecture and how to secure a proverbial infrastructure and how information flows. The question was hinted at earlier, and we need to know this on this side, who is it that you deal with? Department of Homeland Security, Mr. Brecheen brought that up. From my understanding, and this is what I am trying to get clarity on, from my understanding it is there is 3 entities: Department of Justice, Department of Homeland Security, and Department of Defense all touch our communication capabilities above the ground and below the ground. Can you add clarity for me on who you deal with directly? Is there one more than the other?

The discussions I have had with our departments is they kind-of hand the football off, and I really can't find anybody who is running point on this. I will start with you, sir, and we can move back and forth.

Mr. ZERVIGON. I mean, from our experience, I think Customs and Border Protection are showing a lot of leadership on this issue and understanding that this is an architectural problem that needs to be remedied. Obviously with the cost-benefit analysis of being able to do this over a period of time.

Mr. LUTTRELL. Is that brick-and-mortar facilities that our undersea cabling runs into, that, you know, Salt Typhoon is having a heyday with, things like that?

Mr. ZERVIGON. All of them. All the above. So it's about any network connection, any network endpoint that needs to be updated for post-quantum cryptography.

Mr. LUTTRELL. Mr. Hansen.

Mr. HANSEN. As an example, we in the Chrome Browser back in 2023, changed the implementation of the encryption to begin to be post-quantum crypto-resistant because everyone would use it. Right? It's used broadly in the industry. So our strategy is to, whether it's undersea cables, whether it's data centers, whether it's the hardware, make it secure by default.

Mr. LUTTRELL. Is that your company specifically that is providing the security profile for that or is that something that Homeland is coming in assisting with or Department of Defense is coming in and assisting with? I got to tell you, this was kind-of, and I hate to say, ignorant to really kind-of what the answer is to that.

Mr. HANSEN. Yes. In a world where every one of these departments or, you know, sort-of the scope of their oversight is digital or increasingly digital, we work across all of those entities you've mentioned and more on these kinds things.

Mr. LUTTRELL. I feel like we are not doing enough. Case in point, Mr. Graham, with what happened with Claude, and you guys have Gemini, correct? Am I saying that correctly?

Mr. GRAHAM. That's right.

Mr. LUTTRELL. Where the bad actors, the nefarious actors, are utilizing AI capabilities to hack into the kind-of the sweet spot of what we are not looking at.

Mr. GRAHAM, was the—was it a human or software that found the attack or both?

Mr. GRAHAM. On our side it was a combination of both. First, there's a series of detection measures that are generally automated and software-based. This triggered a human investigation that allowed us to—

Mr. LUTTRELL. So as fast as we are moving on the advancements of artificial intelligence and we can't—I don't think we can stop. Because if we slow down, everyone else is going to keep going. Then if we are behind now, we are absolutely going to be in last place. So here we go. If we move to a point where artificial intelligence removes the human element, but you needed the human element to find it, what happens?

Mr. GRAHAM. I am enormously optimistic about the opportunities here to leverage AI to do this. This is the first time we're seeing some of this.

Mr. LUTTRELL. We all are, too. This is us being overly cautious. It is not us that is going to be able to regulate it. It is too fast. By the time you show up in front of us to tell us what happened, whomever took ahold of Claude to make—are they lying in wait? Are they sleeping inside the program now and we have missed it, and they are watching you fix the problem and they know how you fixed it, and they are going to attack someone else that is not as strong and capable or yourself or Google?

Mr. GRAHAM. Well, in this case, it wasn't Anthropic itself that was infiltrated.

Mr. LUTTRELL. Yes, I am sorry. OK.

Mr. GRAHAM. It is very clear that sophisticated actors are now doing preparations for the next time, for the next model, for the next capability they can exploit. This is why we have to be detecting them as fast as possible and mitigating at the model layer.

Mr. LUTTRELL. Because I am going to use the term super scientist. This is what AI has created. You have titrated hundreds of attackers down to 2 or 3 that have the capability to ask the AI the question on exactly how to get in—

Mr. HANSEN. Yes, I think once—

Mr. LUTTRELL [continuing]. At a speed that is uncomprehensible.

Mr. HANSEN. To this point, we've been using behind Gmail and behind the Play Store and behind Chrome for almost a decade AI in its earlier forms to do exactly what you're talking about, so no humans involved. So your question is correct. It's actually been happening, you know, long before the large language models emerged.

Mr. LUTTRELL. OK, thank you.

I am sorry, Mr. Chairman. I yield back.

Mr. OGLES. The gentleman yields back.

I recognize the gentlewoman from New Jersey, Ms. McIver, for 5 minutes.

Ms. McIVER. Thank you, Mr. Chair and Ranking Member, and thank you to our witnesses for joining us today.

Every community, State, and country will be impacted by the benefits and risks of AI. In fact, we already see these impacts occurring. While the United States has been a leader with AI technology, our rivals are innovating in this area with great speed and we have to make sure working people here have what they need to stay safe and successful.

Education will be key to maintaining American dominance, security, and economic success. With my colleagues, Representative Cleaver and Senators Blunt, Rochester, and Hirono and Schiff, we introduced the Workforce of the Future Act. This legislation would help us better examine the skills necessary for workers to thrive in the AI-dominated economy. It will also provide resources for educators and students to get the skills they need to participate in the work force of the future and stay protected against adverse consequences of new technology.

We need to make sure that all Americans are set up to succeed in a world impacted by AI, not be displaced by it. An AI-competent work force will lead to a more secure United States and a stronger future for working people.

With that, Mr. Coates, I would love to talk with you about Trump recently signed an Executive Order that would overturn any State-based AI regulation deemed burdensome. What are some risks of letting AI develop unregulated?

Mr. COATES. I think the important piece with AI regulation is to set clear guidelines and rules of the road and establish transparency amongst the creators. We want to motivate innovation and ensure that the United States stays as a leader in the world on AI.

One of the challenges in cybersecurity in particular can be a patchwork of regulations across States to deal with, especially in things like data disclosure, breach responsiveness, et cetera. So we want to make sure that in the fast-moving field of AI innovation, we are setting the right objectives clear, so we can operate to rules of the road, but we don't hamstring our technology organizations and prevent innovation. The last thing we want to be is on our heels or second to others in the world with AI technology.

Ms. MCIVER. Thank you for that. Just a follow-up, you mentioned cybersecurity. Can you expand a little bit of how important will AI knowledge and competency be in the future of cybersecurity?

Mr. COATES. I would consider AI to be a critical piece of the future of cybersecurity, both from the operators and the defenders. Understanding the core principles of cybersecurity through education, understanding how technology works, and then understanding how the different resources can be used as a defender. As I mentioned in my testimony, there's no question that for defense to be effective, it's going to have to move at the speed of computers. So we need the best humans to understand this technology and harness AI in a defensive capability.

Ms. MCIVER. Thank you for that. As AI data centers continue to expand, how do you balance innovation with the significant environmental and economic burdens they place on local communities and infrastructure?

Mr. Coates, you can start, but anyone else can chime in as well.

Mr. COATES. Maintaining dominance in AI is multifaceted. It's from the technology innovation in the models themselves to having

sufficient power and technology and data centers to fund and power this innovation. So I do think it's critical to work across the Nation to understand where can we have the right locations of data centers with sufficient power. We don't want to lose control of the pieces that go together to build technology. To have effective AI, you have to have sufficient power and data center resources.

Ms. McIVER. Thank you. Anyone else? Mr. Hansen.

Mr. HANSEN. I was just going to say, yes, I talked a little bit about my son's situation and the science and tech and you think of this Alpha Fold, which was the protein folding work that won the Nobel Prize from Google last year. Fusion and energy and clean and safe energy, for me, is another problem. Like the cobbler's children, let's use the AI to help solve that problem. You asked a very good question and that's why we need to keep going on the science and technology as well.

Ms. McIVER. Got it. Anyone else in 20 seconds? All right. Well, thank you so much.

With that, Mr. Chairman, I yield back.

Mr. OGLES. The gentlewoman yields back.

You know, appreciate the topic she touched on because, you know, as we move forward, and hopefully we will have time to come back to it, but this idea of what does that regulatory landscape look like and, you know, this ever-developing, quickly-evolving subject matter where energy is a factor, right? You know, this latency period where we are realizing we have these vulnerabilities that we are not quite ready to, you know, adapt to or backfill. So this is one of those—again, this hearing is the beginning of a very large conversation, whether it is energy, whether it is homeland security, and, quite frankly, the future of our role in the world.

I recognize the gentleman for Alabama for his 5 minutes of questions, Mr. Strong.

Mr. STRONG. Thank you, Mr. Chairman, Ranking Member. Witnesses, thank you for being here today.

Dr. Graham, as my colleagues have mentioned, one concern is that AI allows adversaries to scale operations without scaling personnel. This changes the threat calculus for the United States. When AI tools are misused by cyber activity what visibility, if any, does DHS and CISA have into these incidents?

Mr. GRAHAM. While I'm not familiar with the specific visibility of DHS and CISA here, I do know that what's important is industry should have information-sharing mechanisms with Government in these areas in order to give that visibility and also, in reverse, to understand the areas that industry should defend.

Mr. STRONG. Absolutely. Turning to you, Mr. Hansen, cloud platforms now underpin Federal networks, critical infrastructure, and, increasingly, AI enables Government systems. From a national security perspective, does that concentration of sensitive activity in the cloud create new, wide-spread risk for the homeland?

Mr. HANSEN. Actually, I think it is helping us clean up legacy technology issues. When you look at the vulnerabilities we've had over the last, you know, decade, it's generally people running on old versions of software that they're not maintaining. So we need competition in the space and I think it is competitive in many di-

mensions. But overall, modernizing is going to make you more secure in the moment.

Mr. STRONG. I agree with you. Competition is where it is going to be, also.

AI and data centers are the future. I represent a State that is blessed with all forms of energy: coal, hydro, gas, solar, and nuclear power. We are able to meet the demand. What are your thoughts on AI and data centers in the future?

Mr. HANSEN. You know, I know there's a—this is a big topic, as you would imagine, at Google, and there may be better, you know, people to talk about it. I would just say to the point about using AI, we use AI in the management of our data centers, in the management of the power in a variety of ways. So using the technology to help us do it as efficiently and effectively as possible is sort-of my only perspective. But we could go deeper on that with others in the company.

Mr. STRONG. I also know that companies like Google, Meta, which both of those are located in my district, work closely with universities and the public sector on emerging technologies. In my district, we have institutions such as the Alabama School of Cyber Technology and Engineering that focuses on building early hands-on cyber and technology skills.

Mr. Hansen, from your view, how can public-private partnerships and collaboration with universities help accelerate practical understanding and to secure adoption of AI and cloud technologies across the Government?

Mr. HANSEN. It's a really great question and relates to the work force question as well. We, in fact, over the last few years have stood up what we call cyber clinics. These are not just with the big State universities or private universities. They're with community colleges and they represent places across the country. So I think the working together on the curriculum, the technology, the approach for the next generation is critical.

Mr. STRONG. Thank you. Mr. Zervigon, many national security data sets must remain secure for decades. What are the biggest practical challenges to deploying quantum-resistant encryption at scale today?

Mr. ZERVIGON. The desire to do so, I think. I think the capabilities are there. There are many innovative technologies and innovative companies that can assist. With the desire to do so, I think we can start going by protecting the transport layer, right? The overriding layer, which this information, this data travels.

Mr. STRONG. Thank you. How can Government and industry work together to reduce risk without disrupting operations or slowing innovation?

Mr. ZERVIGON. Looking at it from an architectural standpoint, it's not just about the math. It's not just about creating new algorithms. It's about creating an architecture that allow you to deliver these algorithms, be able to swap them out at scale, be able to protect ourselves in the case that an algorithm is broken, because it will happen. So by doing so, it allows us to mitigate the effects, the ill effects of a harvest now, decrypt later attack.

Mr. STRONG. Thank you. To close out, I would like to ask all the witnesses, if resources are limited, what should DHS and CISA

prioritize first to reduce cyber risk most effectively? I will start on the end.

Mr. GRAHAM. I think establishing threat intelligence-sharing channels, very important. Identifying infrastructure that needs to be secured, that we can go secure.

Mr. STRONG. Thank you. Mr. Hansen.

Mr. HANSEN. Modernization. Right? This is not something we go backward on. We got to go forwards.

Mr. ZERVIGON. Again, looking at the transport layer, looking at the biggest pipes carrying the most important pertinent data, and protect those first and then move downward from there.

Mr. COATES. It would be information sharing on emerging threats and adoption of autonomous defense systems.

Mr. STRONG. Thank you. Mr. Chairman, I yield back.

Mr. OGLES. The gentleman yields back.

I now recognize the gentleman from Louisiana, Mr. Carter, for his 5 minutes.

Mr. CARTER. Thank you, Mr. Chairman.

Cybersecurity is no longer a hypothetical risk. It is a real and growing threat to Louisiana and to our Nation's energy security. Louisiana sits at the heart of America's energy system, with refineries, petrochemical plants, pipelines, LNG export terminals, offshore platforms, and the electric grid all tightly interconnected. A successful cyber attack on any one of these systems could ripple across our entire national economy.

In 2021, the Colonial Pipeline cyber attack shut down a major fuel artery, caused shortages across the Southeast, and drove panic buying and price spikes, all without a single physical asset being damaged. That attack showed just how vulnerable our energy systems can be. That is why we must act now by strengthening cybersecurity, modernizing systems, sharing threat intelligence, and using AI defensively to stop attacks before they succeed.

Mr. Coates, in your testimony you state that bias in AI systems, whether intentional or unintentional, can affect how software is generated, how alerts are prioritized, how decisions are made. How can bias enter AI-driven security tools and what risk that poses to our cybersecurity?

Mr. COATES. It's an excellent question. The challenge in front of us is that we are off-loading decision making into AI when we use AI in our software systems. AI itself is trained on pre-training data, post-training data, configuration, et cetera, but that's reflective of the entity and organization that creates it.

CrowdStrike just released a report recently showing that the DeepSeek LLM model has bias. When you ask that model to create software and mention terms related to items like Tibet and other things not favorable in the CCPI, it generates code that is more vulnerable than had you not mentioned it. So this bias is built deeply into it. Maybe that is unintentional and a result of training data that was used. But nonetheless, we need to be aware that if American corporations are using software that's powered by LLMs, that are built outside the United States, that bias could come back to put us in a more risky position.

Mr. CARTER. So what should we, should the Federal Government, should Congress, be doing to detect and mitigate these actions going forward?

Mr. COATES. The most important piece here is transparency. Requiring in the bill of materials for software procurement that we clearly state the origin of the pieces of the software. This is something we're doing already, but needs to be expanded to cover things like LLM, including where it was created, training information, et cetera.

Mr. CARTER. Dr. Graham, you predict these attacks will only grow in effectiveness. What steps should we be taking to get ahead of this evolving threats, particularly those targeting critical infrastructure? What should Congress, what should we be doing as this committee do, to arm you, to arm others, to make sure that we are not playing catch-up, but we are catching this before it happens?

Mr. GRAHAM. The very first thing we should do is that industry and Government should share threat intelligence so that we can get ahead.

Mr. CARTER. Is that happening at a rate that you are comfortable?

Mr. GRAHAM. It should always happen faster and more. The second is that I believe Congress can enable the deployment of these tools defensively. We can identify the infrastructure we should proactively defend and we can support or remove barriers to pulling these tools in order to defend them.

Mr. CARTER. Mr. Hansen, as CISA developed and issued AI guidance, it worked in collaboration with our international allies. Why should the United States continue to coordinate with countries in this area?

Mr. HANSEN. I was thinking about this when I was in Poland just after the Russian invasion of Ukraine, and they explained how they were now getting grain on the railroad out of Ukraine through Poland, but it had to be changed at the border because the Soviet-era railroad tracks' gauge was different from that in the West. I view this the same. We want American technology to be the railroad gauge of the 21st Century. So, to me, it's a national security question that people use our technology and not others.

Mr. CARTER. Mr. Zervigon, I've got a lot of good friends in Louisiana with that name, so we will check boxes and see if Luis or some of those people are related to you, but.

Mr. ZERVIGON. They are.

Mr. CARTER. Are they really? Fantastic. Some of my very dear friends.

But now that we have had a family reunion, tell me about investments. Are we making the kind of investments to stay ahead of the nefarious actors? As was mentioned earlier, we know that the bad guys sometimes get a lot more information than we do, and their technology grows pretty quickly. What can we do to make sure—because we have got listening ears here, and this is a great bipartisan group of individuals who really want to help. I know my time has expired, so can you give me a quick answer on that?

Mr. ZERVIGON. Sure. I mean, as I mentioned in my testimony, I think increasing the budget for the migration. Right? I think we don't have to do as much on the inventorying and the assessing

and the understanding. We know the pipes that we need to secure, we know the data that we need to secure. We need to start doing that. Also I think helping that is accelerating the time lines and removing these artificial numbers out in the distance. When we should start doing it now.

Mr. CARTER. Thank you, Mr. Chairman. You are very generous.

Mr. OGLES. The gentleman yields back. Thank you, sir, for your questions.

I am going to go to the gentleman from Texas—

Mr. LUTTRELL. Thank you, Mr. Chairman.

Mr. LUTTRELL [continuing]. Mr. Luttrell, for a second round.

Mr. LUTTRELL. The amount of data centers that we are building out, they draw a lot of power, and we are steadily increasing the footprint of each one of those facilities. Now, Texas stands alone as far as the national grid goes. There will come a time the amount of power drawn on everything that we are putting onto the grid will kill it. I am not talking—I am talking next year, 2 years, maybe max. Then what?

I think because we are all in the game together, is there a way that you all can decrease the amount of power, photon communications, or how the grid—how the data centers themselves communicate instead of that amount of power being drawn in? Because we will never catch you. There is no way we can build out enough infrastructure to power the amount of data centers being built. Just those alone.

So I don't know if this is more of a question than a concern that I am sure you are thinking about this. There is going to come a hinge point that it is either going to be an all-stop evolution we have to deal with. We have to do what we have right now because China, they don't have that problem. They are building hand over fist just to keep up the amount of energy that they are drawing. What do we do?

Mr. HANSEN. So, you talked a little about the fusion or technological investment. So I think that's—we need to get started on doing that. We also—and you've seen this from Google, with our TPUs, which is a different type of chip, there are more efficient ways to do some of the computational work related to AI. So I think we need a round of innovation, which we're investing in, to make these chips more efficient and more performative.

Mr. LUTTRELL. Well, that happened—

Mr. HANSEN. That's the work.

Mr. LUTTRELL [continuing]. Before the grid failed?

Mr. HANSEN. That's the work. Yes, that is the work.

Mr. LUTTRELL. Mr. Graham, Mr. Coates, anything on this? I mean, Ms. McIver, hit the nail on the head here. This is a very real thing and we are not trying to slow innovation in any way, shape, or form. The entire globe is moving to the metaverse and we have to be able to sustain that. We do not have the infrastructure in place. I think in Texas, it is 2 years it is going to hit, and I would bet you a dollar on that one. But anyway, thank you, sir.

I yield back.

Mr. OGLES. The gentleman yields back.

I will go to the gentleman, the Ranking Member from OI&A, Mr. Thanedar.

Mr. THANEDAR. Thank you, Chairman Ogles. Appreciate it.

As cyber attacks evolve, it is critical that the private sector share information about cyber threats with the Federal Government. This evolution is only accelerating due to AI, making it more important than ever that the Federal Government has the information necessary to understand current threat landscape. The Cybersecurity Information Sharing Act of 2015, the law that facilitates this kind of critical information sharing between the private sector and Federal Government, this law is set to expire on January 30.

My question to all of you is how important is it that Congress pass a long-term reauthorization of CISA 2015, particularly in light of the rapid evolution and deployment of novel technologies?

Mr. COATES. I think this is critical. In cybersecurity defense the basic primitives are known across organizations. We understand the plumbing, the core items that we need to do, but the techniques and the methods being used by the adversaries continues to change. It's crucial that organizations can say we've discovered this piece and share it with others. So collectively, we don't need to compete on defense, but look at it as a national imperative that we are secure and information sharing is a key piece of that.

Mr. THANEDAR. Thank you.

Mr. HANSEN. Yes, we're very supportive. In fact, I go further and say the Information Sharing and Analysis Centers, the ISACs, which exist by sector, this isn't just going to be a technical issue. This will be a health care, energy, and so the sector-specific sharing we need to focus on as well, particularly as AI operates more at the human layer than at the technical layer.

Mr. THANEDAR. The private sector is usually on the top of the developments and certainly would be in a position to help the Federal Government, right?

Mr. HANSEN. Absolutely. One of the reasons I came to Google from after working in financial services for many years was the realization that everyone was going to—every industry would need the benefits of security being baked into their technology, which includes sharing and making it easier for people to defend themselves.

Mr. THANEDAR. Thank you, I appreciate it. I yield back.

Mr. OGLES. The gentleman yields back.

You know, there is a lot to unpack here and so we will drop in—unless other Members come in, we can drop some the formality and have more of a conversation and feel free to jump in.

You know, I guess I want to start us off with, is we know that we have a lot of, I think, infrastructure gaps. I mean, you know, I like to say we are the dominant predator currently across landscapes, but in this space in particular, that can change rapidly. So when you are setting the marker down, if you had to predict, and whoever wants to answer and understanding this is just a prediction, you know, when you think of our nearest adversary, how long before they are at quantum computing? I know that is a big question by the way, but who wants to guess?

Mr. ZERVIGON. That would be the \$64,000 question.

Mr. OGLES. Right. But are we talking about 2 years or 12 years?

Mr. ZERVIGON. Well, I think the better analysis is whatever the number is, the data that you want to keep secret and you want to

keep protected, is it outside of that? So if you think that a quantum or cryptographically-relevant quantum computer is 5 years out, then any information outside of the 5 we know is problematic. So we need to make sure that we're protected. It's not like Y2K with one moment in time where we need to worry about. It's that moment in time and then the predating of that information and protecting that information.

Mr. OGLES. Well, that is kind-of where I wanted to take this, is that when I think about, you know, just in general, we as individuals, Members of Congress, you know, kind-of device hygiene, the amount of information that is stored that if compromised, that is suddenly is unlocked or unleashed. My fear is currently, as has been stated, is there is a harvesting going on of information across sectors.

So, you know, financial services, that actually is what piqued my interest in AI was being on the Financial Services Committee and specifically the Subcommittee on National Security. I am thinking about all of the threats and how they are escalating and continuing to escalate when it comes to personal information, but also breaching of accounts where suddenly your voice, if it is out there somewhere, can be replicated, where, you know, IDs can be falsified, et cetera.

So, you know, if you want to speak to the amount of information and then what do we do with it? Like how do we—do we need to take this information off-line? Do we silo it? How do we clean up this mess, all these footprints and fingerprints that we have all left across that cyber landscape because it is being harvested, quite frankly, to be weaponized against us?

You want to start, Dr. Graham?

Mr. GRAHAM. I think there are a number of very substantial opportunities that we have here. I'm, again, I'm extremely optimistic about using AI to help do this. Anthropic takes privacy and the sensitivity of data extremely seriously. I think we could probably unleash quite a lot of innovation here using AI to secure data infrastructure sensitive systems. I think this is going to be one of the important topics if we deploy this technology more and more into the economy to ensure that it's critical we get it to defend critical infrastructure without exposing it anymore.

Mr. HANSEN. Yes. First of all, the reason we implemented the new encryption in Chrome was to start to get ahead of exactly the kind of question you're talking about. So there are some common utilities, whereas we at Google or other companies migrate, you get an architectural benefit for others.

But to the point on using AI, we have used, again, even before large language models, AI to help identify unused data, label data per certain sensitivities, and then you can implement policy that protects it. But I think, you know, he's correct. We'll have to use AI to get to the scale of the problem that you're describing. That means we'll also have to modernize, though, because we can't do that with the servers that are under desks and in, you know, sort-of second-class data centers that no one's modernized before. So that combination of modernization and using the tools, I do think we can scale to that problem.

Mr. COATES. I see two parts to the question you raise, one of which is how do we defend organizations against the rising orchestration of attacks that we've talked about some through AI? The second piece around how quantum changes things, and the biggest challenge with decrypting—the ability to decrypt traffic when quantum becomes relevant is the change that we need to do to be defensive here is a administrative and operational change.

We understand the systems that we have inside our organizations. We need to essentially upgrade them. Unfortunately, with the number of priorities we have for cybersecurity, it needs to become a top issue for organizations to say this needs to happen by this date, because otherwise, we're going to be really caught behind the eight ball where the data will be captured, it will be decrypted, and the time to do the upgrade will be so significant that we'll be in that risky position for a much longer period.

Mr. OGLES. Thank you. You know, Google's infrastructure, you mean, the amount of computing that you are supporting, from Government to private to health. I mean, just across the board, when you look at these kind-of constant attacks, so just had a hearing last week, Financial Services, on the Oversight Committee. We had, you know, everyone from Verizon to, you know, the credit card companies to, you know, across the board. Right? The social media platforms, the architecture platforms. We were talking about the threats that they are facing and the amount of investment that is being made and, quite frankly, leveraging.

So when it comes to credit cards, for example, it is where you have AI that is constantly watching transactions, looking for those patterns that otherwise are outside the norms. But what are those fail points when you look at that ecosystem from a Google perspective?

Mr. HANSEN. Yes, it's a great point. I'll maybe just extend that a little bit and see if this is what you're asking about. But it is the controls that we care about in finance or health care or transportation are going to be different, the risks are different. So it's not just about the plumbing, let's call it, the technology, but in your credit card, the limits you set. Show me what—any transaction over \$100 and you get that monitoring. You think about the kind of monitoring that occurs in health care.

I think the key is that this isn't just a technical problem. This is an industry problem. AI can help because AI understands the language. If you write a policy that says this heartbeat level is problematic under these conditions, the AI model is going to be better at monitoring that than a human. So that's where we need to go, is to use AI. This is my—I keep coming back to the cobbler's children. Let's not, you know, be, you know, shoeless in defending ourselves.

Mr. OGLES. Well, again, on the AI, you know, when I think about—when you look at Elon and some of the other companies that are doing the—any of the autonomous robots or humanoids, whatever you want to call them, and the ability to have a partner that now can watch a child who is ill or a spouse or an elderly parent that is—where they are wearing a ring or a bracelet, where they are constantly being monitored in real time, where you have a situation where they can dispense or disperse medicines and,

again, immediately relaying back to the doctor, there is a huge upside to this. It is going to be transformative in a way that, again, I think is hard to fathom.

My concern is when we have these nation-states that are constantly seeking to exploit what otherwise could be used for tremendous good. So I do think when I think about China and their overt—I mean, at this point, they are not even hiding it. I mean, you know, I think they were testing. You know, the question or the point was made is, you know, I don't think we should ever underestimate our adversaries. This idea that, you know, they put it out there, it was detected, you know, they are watching to see how you detected it. How can they replicate or do it better the next time?

So we know it is coming, it is just a matter of time. You know, as we think about—and the investment, quite frankly, that they are making is that, I think, you know, from our perspective, we have to do a better job. You know, put up the guardrails, increase the transparency. But this flow of information is going to be critical. That is going to include some of our partners overseas. So from an industry perspective, how is that cross-collaboration going with some of our European partners or Israel or to the extent that you can disclose?

Mr. GRAHAM. On topics of national security, Anthropic works with U.S. and democratic allies quite heavily for exactly this reason. One of the areas of collaboration that has helped the most has been in testing of model capabilities, so that everybody understands where we're at and what's coming down the pipeline. That is the key first step.

Additionally, there are probably international insights into, how we do secure our infrastructure and learn from each other? Broadly, we generally support this, and I think it's a testament to America's leadership that it has instigated that degree of international collaboration.

Mr. HANSEN. It's a great point. Just my job's changed dramatically from the, you know, 20 years ago when I started. I was just thinking this year, I was in Tokyo, Singapore, Abu Dhabi, Tel Aviv, Sao Paulo, Warsaw, talking exactly about these kinds of issues and how do we raise the baseline for those citizens? So it's a big part of the job. We realize that.

Mr. ZERVIGON. For us, I think a large part of it is on the architecture, right? As we develop the architecture that allows different countries, different regions to employ the encryption that they want to employ, we certainly like to show leadership in that. We are with the work that NIST has done over the past decade. But at the end of the day, different countries, different regions are going to want to do what they want to do. So focusing on the architecture enables that.

Mr. COATES. In terms of information sharing I would point to the innovation pipeline. I was just in Tel Aviv last week at a major cybersecurity conference, speaking with start-ups and other innovators in the space. Tel Aviv in particular and Israel creates amazing technology that bridges to the United States as one of their main customer bases.

So as we look at where the next great ideas are coming from, they are being created inside the United States and they're being

created with our allies. Working closely, especially with Israel, for cybersecurity is definitely to our advantage.

Mr. OGLES. Well, on that, when I think about the innovation and the innovation pipeline, you know, as we look at the NFI 7, kind-of 14 Eye groups, you know, I think one of where it is imperative that we are sharing information across kind-of countries and nation-states is this, you know, certain countries based off of where they are at and the type of threats they are exposed to get quite good at those types of attacks. So what South Korea is facing may be slightly different or a different perspective than Israel is facing versus Eastern Europe.

So one of the things that I have done is I have had the opportunity to travel in South and Central America and to Eastern Europe to talk about cybersecurity. What troubles me is in many of these countries, especially when you get into that second tier, is they are wholly unprepared.

I think, Mr. Zervigon, you mentioned that, you know, what we want to do is create a cyber environment where the world is, quite frankly, reliant on our architecture, our expertise. So the idea of the chips, there's some huge—you know, it is a pause moment to figure out what do we want to share versus where do we want to hold back. That is probably not a conversation that we can have in this setting.

But that being said is ultimately we want our global partners, whether in South America or Africa or Europe, Central America, to be dependent on us and trust us in this ever-evolving space. Because in my humble opinion, the threat to the West and the developing world is China. It is time we have that honest conversation. Quite frankly, your report really puts a fine point on the fact that this was an intentional attack to undermine the United States of America, to undermine the West and to, quite frankly, to try to achieve a technical advantage that they currently don't have as they seek to leap forward in their own development and their own technology.

So with that, and, you know, we are probably going to end a little soon, but what I would love to do is just go down the line, any thoughts that you might have. You know, sometimes you are in a room, you don't ask the right questions, so feel free to point out the right question. Then also, what is that thing? You know, what are next steps? Then what keeps you up at night?

Dr. Graham, you are at the top of the table, so we will just start with you, sir.

Mr. GRAHAM. To me personally, as we watch these threats, and have for the past 2-plus years, we have seen the models go from zero to extremely useful and now used in the real world. This only happens because we monitor this threat in the first place.

But the most important thing in our team's view from now on is to take this moment here as the change point, is from now on that we will have a degree of scale that I think we've never had before and very possibly very soon, a degree of sophistication. I fear the day we wake up and models are doing things more complicated and sophisticated than the best humans on Earth are able to understand.

The only answer we think over the long term is to make sure that we're using models to keep up and outpace the attackers. We need to give the defenders a permanent advantage. We're going to work really hard to make sure our models can do that. We're going to work really hard to make sure that they're deployed. This is a cross-industry challenge. We have to work with Government on it. This is, we believe, the fundamental issue.

Mr. OGLES. Dr. Hansen.

Mr. HANSEN. Yes, maybe just two things. One, I'm reminded that in 2009, Google was compromised by Chinese threat actors. This goes back over 15 years. It was our—it was a watershed moment at the company and we spoke openly about it. They had attacked 25 companies. It's really where the modern architecture for security was born. You hear about zero trust. This was the company redoing our infrastructure from the ground up to be up to the kind of attacks we now knew were possible.

To the point about AI, I think that's the next phase of this threshold is to put in hands of defenders the tools that will allow them to be successful in ways that we've, frankly, been—the numbers game doesn't work for us right now with all this legacy software. So now is the time to put those tools in the hands of defenders.

Mr. ZERVIGON. I would say also to accelerate the time lines and the budget as we talked about. I mean, 15 years ago, two-factor authentication, nobody had ever heard of it. Now it's everywhere. You can't buy concert tickets without two-factor authentication. Same thing is going to be the case with encryption.

I think under the Legislative branch as well as the Executive branch, continuing to lead on this and to kind-of push the envelope and set the table for innovative technologies and innovative companies to actually be able to start doing what they do best rather than waiting for legacy time lines to take hold, I think that's in everyone's best interest. It starts with the Government and then it'll move quickly to critical infrastructure or critical industries and then it'll move to everything, just like two-factor authentication did.

Mr. COATES. The country that leads in AI will lead in the world. This is the most important and innovative time in recent history. I believe that it is imperative that we align behind the challenges may be that data centers, be that energy, be that human resources, be that regulation, to create a transparent playing field in the United States where we can spur innovation forward. I think if we are caught up in any of the obstacles in pursuit of that, it will only give foreign adversaries the upper hand and then let them lead other countries to build on top of their technologies, which will be even harder to dig out from.

So the future is in front of us and leading in AI is the most important thing we can do.

Mr. OGLES. Absolutely. You know, I thank all the witnesses. Mr. Coates, to your point, you know, of there are a lot of subjects in Congress that we address that are kind-of very heated and at times partisan, but I would like to think this is the one that isn't. We have a lot to do, whether it is the sharing of information, whether

it is better educating our allies overseas, preparing for that—the energy load that we know is coming, and just sheer innovation.

Like has been said, you know, we want to put up the guardrails to protect Americans and our allies. We also understand that our adversaries are not going to use guardrails. I would argue that they would quite—they, quite frankly, are willing to be reckless in achieving this goal, this endgame, which is AI and quantum. Because it does, it changes the world forever.

So I think this is the wake-up call. This is that moment in time that we will point to in this space. Did we heed the warning? Were we listening? Were we paying attention?

You have got our attention. My challenge to you would be to feel free to come to this body, come to me, come to the Ranking Member and have those honest conversations of we see a deficiency here and we need your help. Or this is a space where you are getting it wrong. Because if we don't have that communication and that trust, forget ideologies and politics and who you voted for, this is about national security. This is about your son. Right? Is not putting impediments and guardrails in the way that impedes that cure or whatever discovery is next. I truly—I can't imagine what the future looks like, but it's coming whether we prepare for it or not.

So I commend all of you for being here. Quite frankly, I would love to have the conversation with each of you about having a working group that is outside that reports back to this body. We can get bipartisan membership to participate in it, so to guarantee that we truly—is it one of the things to get platitudes, right? It is one thing, oh, we are going to share information, we are going to work with our allies. We are going to do the right thing for the right reasons. But if we are not having any conversations, it is all platitudes. I am not one to shy and beat around the bush. If we don't get this right, we are screwed. Right?

I think you said, Dr. Hansen, you know, the defender has to be right every time. Right? Your adversary only has to be right once. If we mess this up, it changes everything forever.

Any final thoughts?

Well, I, again, I thank you all. I am humbled that you would come before Congress. It is important that we have this conversation. I look forward to getting to know each of you better. I personally will reach out to each one of you individually, so that you know that you have access to Congress every single day of the week, 24/7. I will answer my phone.

With that, the committees stand adjourned. God bless you, sir, and your son.

[Whereupon, at 12:35 p.m., the subcommittees were adjourned.]

APPENDIX

QUESTION FROM HONORABLE JAMES R. WALKINSHAW FOR LOGAN GRAHAM

Question. What are your recommendations to ensure that safety and security of artificial intelligence (AI) models scale and extend beyond how we think about model development in today's graphic processing unit era and into a far broader landscape brought about by quantum computing and quantum machine learning?

Answer. At Anthropic, our work on the Frontier Red Team is premised on the idea that safety and security measures must be built proactively and evaluated continuously. Three recommendations from my testimony are directly applicable to ensuring that foundation holds as compute architectures evolve.

First, codify and expand model testing capacity. The U.S. Center for AI Standards and Innovation (CAISI) has developed real expertise in evaluating frontier AI models for national security-relevant capabilities. Congress should permanently authorize CAISI and resource it to develop evaluation methodologies that can adapt to new capabilities over time. The voluntary agreement Anthropic has with CAISI provides a replicable model for how this can work in practice.

Second, strengthen threat intelligence sharing between frontier AI labs and the U.S. Government. The CCP-backed campaign we disclosed demonstrates that threat actors are already probing frontier AI models to leverage their capabilities for offensive cyber capabilities and other malicious use cases. As those models grow more capable, the imperative for robust, real-time intelligence sharing between Government and industry only increases. Congress should establish formal channels modeled on existing critical infrastructure information-sharing mechanisms.

Third, maintain and strengthen export controls on advanced compute. The strategic logic here extends to any hardware paradigm that could provide adversaries with the capacity to develop or run frontier AI systems. Ensuring that authoritarian nations cannot acquire the advanced compute needed to close the gap with U.S. frontier capabilities is the single most important structural safeguard we have.

The most important thing Congress can do to ensure AI safety and security scales into the future is to build the institutional infrastructure—testing capacity, intelligence sharing, and compute controls—that can keep pace with a rapidly-changing landscape.

QUESTIONS FROM HONORABLE JAMES R. WALKINSHAW FOR ROYAL HANSEN

Question 1. How can digital transformation and transitioning to cloud computing support an organization's cybersecurity objectives?

Answer. Google keeps more people safe on-line than anyone else, and this scale has required us to deliver pioneering approaches to cloud-native security. As a result, Google Cloud defends its users' data against threats and fraudulent activity using the same infrastructure and security services it relies on for its own operations.

With respect to this infrastructure, Google Cloud provides a secure-by-design foundation—a model for risk management supported by products, services, frameworks, best practices, controls, and capabilities—that acts as an organization's security transformation partner. By building advanced security into every stage of our product development and cloud infrastructure, we enable organizations to modernize and strengthen their IT security while helping users protect their personal information and access the internet safely.

As referenced below, Google Cloud enables organizations to implement a zero-trust approach—where trust in users and resources is established via multiple mechanisms and verified on a continuous basis—to protect their workforce and workloads.

Question 2. How can the scale of cloud computing assist with mitigating cyber events?

Answer. Google Cloud’s baseline security architecture adheres to Zero Trust principles—the idea that every network, device, person, and service is not trusted until it proves itself. It also relies on defense in depth, with multiple layers of controls and capabilities to protect against the impact of configuration errors and attacks.

Public clouds have the scale to implement levels of security and resilience that few organizations have previously constructed. At Google, we run a global network, and we build our own systems, networks, storage, and software stacks. We equip this network with a high level of default security; our Titan security chips assure a secure boot; we provide default data-in-transit and data-at-rest encryption; and we make available confidential computing nodes that encrypt data even while it is in use.

We prioritize security by design and have a team of security engineers who work continuously to deliver secure products and customer controls. Our global public cloud enables Google to achieve unparalleled economies of scale, making security more efficient and cost effective for Google and its customers or users.

Question 3. How can cloud service providers contribute to the secure development of Artificial Intelligence?

Answer. Enterprises today face the critical challenge of delivering AI to production while ensuring accuracy, safety, and data security. Google’s approach to generative AI prioritizes enterprise readiness with built-in mechanisms for robust data governance, privacy controls, IP indemnification, and responsible AI practices. We provide the tools and services necessary to secure AI and offer data sovereignty options, giving customers the confidence to deploy models at scale.

Google Cloud takes several steps to help organizations leverage the power of generative AI:

- *Conducting comprehensive reviews during AI product development.*—Google Cloud identifies and assesses potential risks at both the model level and the point of their integration into a product or service. Our approach considers how AI will interact with the world and existing systems and evaluates the potential impacts and risks that may be posed both at the initial release and at points thereafter. Reviewers understand that potential risks and impacts might be different at the model level and at the application level and consider mitigations accordingly. We draw from various sources, including academic literature, external and internal expertise, and our in-house ethics and safety research.
- *Privately releasing models.*—The private release of models allows our product teams to gather valuable feedback before we make these models generally available. Once feedback is incorporated, we update our product documentation to account for any changes.

QUESTION FROM HONORABLE JAMES R. WALKINSHAW FOR EDDY ZERVIGON

Question. Your company is supporting efforts to protect enterprise infrastructure against brute force quantum attacks on encryption that could cripple e-commerce, personal communication, and national security. Other than post-cryptography standards that the National Institute of Standards and Technology approved in 2024, what other assistance could the Federal Government provide to prioritize or raise awareness of dangers of quantum attacks on our commercial communications infrastructure?

Answer. As requested by the House Committee on Homeland Security, I’m responding to your letter dated February 12, 2026, asking for additional insights on what other assistance, beyond the adoption of post-quantum cryptography (“PQC”) standards, that the Federal Government can provide to prioritize or raise awareness of the dangers of quantum attacks on our commercial communications infrastructure.

As I testified at the joint hearing titled, “The Quantum, AI, and Cloud Landscape: Examining Opportunities, Vulnerabilities, and the Future of Cybersecurity” on Wednesday, December 17, 2025, the most important initiatives that this committee and Congress can undertake are to approve funding for PQC migration now and to accelerate the time lines for adoption on our most sensitive data networks. We cannot achieve innovative results on legacy time lines, and we can’t afford to wait. Congress should work with Federal agencies to accelerate this migration through legislation and regulation.

Two additional areas of concern have surfaced during our work with Federal agencies to deploy PQC’s. These are inter-vendor compatibility and crypto-agility.

As standards are adopted and agencies and enterprises begin to migrate to PQC, we need to ensure that proprietary vendor implementations of PQC’s do not slow our ability to scale. This committee should provide guidance that mandates interoperability between different vendor platforms in their implementation of PQC’s.

Finally, our ability to change PQC's (either the algorithm or the implementation) needs to be as seamless as possible to prevent any delays in adoption of these changes when (because it's going to happen) an algorithm or implementation is broken. We need to be sure that any implementations can support future NIST PQC algorithms, irrespective of legacy technical limitations (such as key-size, packet size, or network quality, etc.).

We very much look forward to our continued work with the committee's staff and welcome any opportunities to offer our expertise around this issue. I thank you all for your leadership on this critical issue and your efforts to strengthen our Nation's security and expand economic prosperity.

QUESTIONS FROM HONORABLE JAMES R. WALKINSHAW FOR MICHAEL COATES

Question 1. How do you foresee the economic development opportunities for advances in quantum computing and how it will shape the cybersecurity and AI markets?

Answer. Quantum computing represents both a significant economic opportunity and a moment of critical security transformation.

In the short term, the most immediate opportunities from quantum computing will not center on cybersecurity risk, but on its computational power to solve previously intractable problems. Quantum acceleration has the potential to impact logistics optimization, pharmaceutical discovery, advanced materials science, energy systems, and manufacturing. These advances could materially increase productivity across multiple sectors. Quantum techniques may also meaningfully enhance certain artificial intelligence workloads over time, further accelerating AI-driven innovation.

At the same time, quantum computing introduces structural implications for cybersecurity. Public-key cryptography underpins nearly every secure digital system, including financial transactions, identity infrastructure, cloud workloads, software updates, and government communications. The transition to post-quantum cryptography (PQC) is not a routine software update. It is a multi-year infrastructure migration affecting hardware, firmware, cryptographic protocols, certificate management systems, and embedded technologies.

This transition represents a substantial economic activity in its own right. It will require software modernization across both public and private systems, along with operational oversight, planning, validation, and testing. Organizations must inventory cryptographic assets, implement crypto-agility, update long-lived systems, and ensure interoperability. The scale of this effort will create significant market opportunities in cybersecurity, infrastructure management, and enterprise modernization.

In short, quantum computing will drive economic development through both innovation expansion and necessary security modernization. The organizations that succeed will be those that enable practical, secure transition rather than simply theoretical advancement.

Question 2. How can the Federal Government ensure that citizens broadly benefit from rapid advances in quantum computing, like they did with the advent of personal computing in the 1980's and the internet in the 1990's?

Answer. Broad economic benefit depends on open standards, distributed innovation, and trusted digital infrastructure.

First, the Federal Government should accelerate adoption of post-quantum cryptography within the public sector and use its procurement authority to drive timely migration across critical industries. Federal systems process sensitive citizen data and underpin national infrastructure. Leading by example reduces systemic risk. Clear time lines and enforcement mechanisms will also push the private sector to modernize more quickly. That acceleration is in the public's interest. Citizens depend on banks, health care providers, utilities, cloud platforms, and other businesses to safeguard their data and operations. A delayed transition increases collective vulnerability.

Second, policy makers should preserve an open innovation ecosystem. The economic success of the personal computing and internet revolutions stemmed from broad participation across start-ups, universities, and private industry. Encouraging domestic research commercialization and supporting start-up formation will help ensure quantum capability is not overly concentrated and that its economic benefits are widely distributed.

Third, work force development is essential. Migrating national infrastructure to quantum-safe systems will require engineers and security professionals trained in both legacy and next-generation cryptography. Without sufficient technical talent, modernization efforts stall and risk persists.

Finally, trust and privacy must remain central. Citizens only benefit from technological revolutions when they trust the systems that underpin commerce, commu-

nication, health care, and financial services. Strong, uncompromised encryption is foundational to that trust. History has demonstrated that deliberately weakening encryption—even with limited intent—introduces systemic vulnerabilities that adversaries can exploit. As the Nation transitions to quantum-resistant systems, preserving robust, trustworthy encryption protects individual privacy, economic stability, and national security.

Quantum's promise will be realized not merely through innovation, but through secure, timely, and broadly-deployed implementation that maintains public confidence in the digital ecosystem.

