

UNLOCKING THE NEXT GENERATION OF AI
IN THE U.S. FINANCIAL SYSTEM
FOR CONSUMERS, BUSINESSES,
AND COMPETITIVENESS

HEARING
BEFORE THE
SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL
TECHNOLOGY, AND ARTIFICIAL INTELLIGENCE
OF THE
COMMITTEE ON FINANCIAL SERVICES
U.S. HOUSE OF REPRESENTATIVES

ONE HUNDRED NINETEENTH CONGRESS

FIRST SESSION

SEPTEMBER 18, 2025

Serial No. 119–43

Printed for the use of the Committee on Financial Services



www.govinfo.gov

U.S. GOVERNMENT PUBLISHING OFFICE

63–432 PDF

WASHINGTON : 2026

HOUSE COMMITTEE ON FINANCIAL SERVICES

FRENCH HILL, Arkansas, *Chairman*

BILL HUIZENGA, Michigan, <i>Vice Chairman</i>	MAXINE WATERS, California, <i>Ranking Member</i>
FRANK D. LUCAS, Oklahoma	SYLVIA R. GARCIA, Texas, <i>Vice Ranking Member</i>
PETE SESSIONS, Texas	NYDIA M. VELÁZQUEZ, New York
ANN WAGNER, Missouri	BRAD SHERMAN, California
ANDY BARR, Kentucky	GREGORY W. MEEKS, New York
ROGER WILLIAMS, Texas	DAVID SCOTT, Georgia
TOM EMMER, Minnesota	STEPHEN F. LYNCH, Massachusetts
BARRY LOUDERMILK, Georgia	AL GREEN, Texas
WARREN DAVIDSON, Ohio	EMANUEL CLEAVER, Missouri
JOHN W. ROSE, Tennessee	JAMES A. HIMES, Connecticut
BRYAN STEIL, Wisconsin	BILL FOSTER, Illinois
WILLIAM R. TIMMONS, IV, South Carolina	JOYCE BEATTY, Ohio
MARLIN STUTZMAN, Indiana	JUAN VARGAS, California
RALPH NORMAN, South Carolina	JOSH GOTTHEIMER, New Jersey
DANIEL MEUSER, Pennsylvania	VICENTE GONZALEZ, Texas
YOUNG KIM, California	SEAN CASTEN, Illinois
BYRON DONALDS, Florida	AYANNA PRESSLEY, Massachusetts
ANDREW R. GARBARINO, New York	RASHIDA TLAIB, Michigan
SCOTT FITZGERALD, Wisconsin	RITCHIE TORRES, New York
MIKE FLOOD, Nebraska	NIKEMA WILLIAMS, Georgia
MICHAEL LAWLER, New York	BRITTANY PETTERSEN, Colorado
MONICA DE LA CRUZ, Texas	CLEO FIELDS, Louisiana
ANDREW OGLES, Tennessee	JANELLE BYNUM, Oregon
ZACHARY NUNN, Iowa	SAM LICCARDO, California
LISA McCLAIN, Michigan	
MARIA SALAZAR, Florida	
TROY DOWNING, Montana	
MIKE HARIDOPOLOS, Florida	
TIM MOORE, North Carolina	

Ben Johnson, *Staff Director*

SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL TECHNOLOGY, AND ARTIFICIAL INTELLIGENCE

BRYAN STEIL, Wisconsin, *Chairman*

TOM EMMER, Minnesota, <i>Vice Chairman</i>	STEPHEN F. LYNCH, Massachusetts, <i>Ranking Member</i>
BILL HUIZENGA, Michigan	BRAD SHERMAN, California
WARREN DAVIDSON, Ohio	BILL FOSTER, Illinois
JOHN W. ROSE, Tennessee	JOSH GOTTHEIMER, New Jersey
WILLIAM R. TIMMONS, IV, South Carolina	AYANNA PRESSLEY, Massachusetts
MARLIN STUTZMAN, Indiana	RITCHIE TORRES, New York
BYRON DONALDS, Florida	SYLVIA R. GARCIA, Texas
ZACHARY NUNN, Iowa	BRITTANY PETTERSEN, Colorado
TROY DOWNING, Montana	SAM LICCARDO, California
MIKE HARIDOPOLOS, Florida	
TIM MOORE, North Carolina	

C O N T E N T S

Thursday, September 18, 2025

OPENING STATEMENTS

	Page
Hon. Bryan Steil, Chairman of the Subcommittee on Digital Assets, Financial Technology and Inclusion, a U.S. Representative from Wisconsin	1
Hon. Stephen Lynch, Ranking Member of the Subcommittee on Digital Assets, Financial Technology and Inclusion, a U.S. Representative from Massachusetts	2

STATEMENTS

Hon. French Hill, Chairman of the Committee on Financial Services, a U.S. Representative from Arkansas	4
--	---

WITNESSES

Mr. Christian Lau, Co-Founder and Chief Product Officer, Dynamo AI	4
Prepared Statement	7
Dr. David Cox, Vice President, AI Models, IBM Director, MIT-IBM Watson AI Lab	18
Prepared Statement	20
Mr. Matthew Reisman, Director, Privacy and Data Policy, Center For Information Policy Leadership	26
Prepared Statement	28
Mr. Daniel Gorfine, Founder & CEO, Gattaca Horizons; Former Chief Innovation Officer & Director of LabCFTC	30
Prepared Statement	32
Dr. Nicol Turner Lee, Senior Fellow and Director, Center for Technology Innovation, Brookings Institution	50
Prepared Statement	52

APPENDIX

MATERIALS SUBMITTED FOR THE RECORD

Warren Davidson:	
Op-ed for the Daily Caller	96
Hon. Maxine Waters:	
America's Credit Unions	99
The National Fair Housing Alliance (NFHA)	103
Public Citizen	111

RESPONSES TO QUESTIONS FOR THE RECORD

Written responses to questions for the record from Mr. Christian Lau	113
Written responses to questions for the record from Dr. David Cox	116
Written responses to questions for the record from Mr. Matthew Reisman	118
Written responses to questions for the record from Dr. Nicol Turner Lee	123

LEGISLATION

H.R. 4801, the Unleashing AI Innovation in Financial Services Act	127
---	-----

IV

	Page
H.R. 2152, the Artificial Intelligence Practices, Logistics, Actions, and Nec- essities (PLAN) Act	148
H.R. 1734, the Preventing Deep Fake Scams Act	153

UNLOCKING THE NEXT GENERATION OF AI IN THE U.S. FINANCIAL SYSTEM FOR CONSUMERS, BUSINESSES, AND COMPETITIVENESS

Thursday, September 18, 2025

U.S. HOUSE OF REPRESENTATIVES,
COMMITTEE ON FINANCIAL SERVICES,
Washington, DC.

The subcommittee met, pursuant to notice, at 2:12 p.m., in room 2128, Rayburn House Office Building, Hon. Bryan Steil [chairman of the subcommittee] presiding.

Present: Representatives Steil, Hill, Huizenga, Davidson, Rose, Timmons, Nunn, Downing, Haridopolos, Lynch, Waters, Foster, Pressley, Torres, Garcia, Pettersen, and Liccardo.

Chairman STEIL. The Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence will come to order.

Without objection, the chair is authorized to declare a recess of the committee at any time.

The hearing is titled “Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness.”

Without objection, all members will have 5 legislative days within which to submit additional material to the chair for inclusion in the record.

I now recognize myself for 4 minutes for an opening statement.

OPENING STATEMENT OF HON. BRYAN STEIL, CHAIRMAN OF THE SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL TECHNOLOGY AND INCLUSION, A U.S. REPRESENTATIVE FROM WISCONSIN

Artificial intelligence is rapidly changing industries across the world. Few sectors are more prepared and more impacted than financial services. For decades, our financial institutions have been at the forefront of development and deploying AI technology, from algorithmic trading to machine learning systems used in risk management to combating fraud in better and faster ways.

The rise of generative AI represents the next transformative step, which could bring new efficiencies, new opportunities, and potential new risks to our financial markets. This subcommittee has demonstrated strong leadership in charting a path forward for transformative technologies, including digital assets, most recently through the CLARITY Act and the Guiding and Establishing National Innovation for U.S. Stablecoins (GENIUS) Act.

Today, we turn our attention to artificial intelligence, examining how it is being deployed in financial markets and assessing whether the current regulatory framework is prepared to keep pace.

The United States has long been a hub for financial innovation, and we must ensure our policies support responsible AI adoption, not stifle it, as seen during the Biden-Harris Administration.

Recently, the Trump Administration released its AI Action Plan, which emphasizes American AI leadership across all sectors. We must look to how we can enhance American leadership and competitiveness in financial technology.

It is essential that regulations strike the right balance, fostering innovation while ensuring investor protection and market integrity. Our financial regulators must be cognizant of that technology it uses, and Congress must provide the clarity to encourage responsible development here at home. It is paramount our markets are not left behind in the global race for AI leadership.

We are fortunate to have with us today a panel of esteemed experts who bring a wealth of knowledge and experience in artificial intelligence and its deployment in the financial system. The information we learn today will build upon our prior work and assist the subcommittee in shaping policies that encourage responsible innovation in financial markets while cementing American AI leadership.

I want to thank our witnesses for being with us today, and I look forward to today's discussion.

I will now recognize the ranking member of the subcommittee, Mr. Lynch, for 4 minutes for an opening statement.

OPENING STATEMENT OF HON. STEPHEN LYNCH, RANKING MEMBER OF THE SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL TECHNOLOGY AND INCLUSION, A U.S. REPRESENTATIVE FROM MASSACHUSETTS

Mr. LYNCH. Thank you very much, Mr. Chairman, and thank you for your courtesy. I had a vote on a subpoena in another hearing and that caused my delay but thank you.

I want to thank this panel of witnesses for your willingness to come before the committee and help us with our work.

This hearing continues this committee's oversight work to examine the use of AI in the financial services and housing sectors. As evidenced by the final report issued by the bipartisan working group that Chairman Hill and I co-chaired last year, financial institutions, fintech companies, and even financial regulators are already deploying AI to maximize operational efficiency and achieve cost reduction across their core functions, from personalized customer services and consumer lending to fraud detection and financial crime monitoring.

At the same time, the rapid development of AI-based technologies has introduced serious risks into the financial service space. The Treasury Department, the Federal Reserve, and Consumer Financial Protection Bureau have all expressed concerns. Other financial regulators have repeatedly cautioned that the development of robust and trustworthy artificial intelligence is dependent on our ability to encourage innovation that maximizes

oversight, consumer protection, data privacy, as well as workforce protection and marketplace fairness.

In view of these considerations, I am concerned by the Trump Administration's decision to rescind some of the commonsense Federal directives that sought to advance the safe and responsible development of artificial intelligence. This includes the reversal of an executive order directing all Federal agencies to enforce existing consumer protection laws and to develop additional regulatory safeguards against fraud, unintended bias, discrimination and privacy violations only when absolutely necessary.

President Trump subsequently issued his own AI Action Plan, a strategy that largely reflects the divergent view that consumer protection stands as an impediment to AI innovation, a position with which I strongly disagree.

Regrettably, the AI Action Plan undermines State-level AI regulation by requiring Federal agencies to, quote, consider a State's regulatory climate, close quote, when allocating funding, and to outstrip Federal funds from any State with a regulatory framework that it considers burdensome.

I would note that recent legislation to impose a 10-year moratorium on State AI regulation was soundly defeated in the U.S. Senate by an overwhelming bipartisan vote of 99 to 1.

The AI Action Plan similarly impedes Federal oversight of AI-based systems by directing the Federal Trade Commission to stand down from ongoing investigations that unduly burden AI innovation. It is the fundamental mission of the Federal Trade Commission (FTC) to protect the American public from unfair or deceptive business practices.

President Trump also recently issued an executive order to limit Federal procurement of AI large language models to those that are developed with ideological neutrality, as reported by the independent Brennan Center for Justice.

Compliance with the order will likely require technology companies to degrade model performance and undermine public trust in their technology. Not surprisingly, our committee has received regular reports from AI stakeholders who have put AI development on hold over concerns that their integration of fairness metrics will run afoul of the executive order.

Innovation will also not be well served by the dismantling of the Consumer Financial Protection Bureau (CFPB), a Federal agency that has sought to advance the development of trustworthy AI systems in financial services through consumer protection enforcement actions and regulatory clarity.

In stark contrast to these actions, the work of our bipartisan AI Working Group stems from a genuine commitment by members on both sides of the aisle to foster AI innovation in the financial services industry while also ensuring that regulators are equipped with the authorities and resources necessary.

In closing, to this end I look forward to continuing to work on a bipartisan basis with Chairman Hill, Chairman Steil, Ranking Member Waters and our committee colleagues to advance the U.S. leadership in this important area.

Thank you, Mr. Chairman, for your courtesy once again, and I yield back the balance of my time.

Chairman STEIL. The gentleman yields back.

I now recognize the chairman of the full committee, Mr. Hill, for 1 minute for an opening statement.

STATEMENT OF HON. FRENCH HILL, CHAIRMAN OF THE COMMITTEE ON FINANCIAL SERVICES, A U.S. REPRESENTATIVE FROM ARKANSAS

Chairman HILL. Thank you, Chairman Steil and I appreciate our panel being with us.

AI has, of course, rapidly changed the way Americans live, work, and engage with our financial system. Last Congress, we explored AI practices at financial services firms, regulators, and supervisors. Congressman Bill Foster and I served on Speaker Johnson and Minority Leader Jeffries' Bipartisan congressional AI Task Force. Ms. Garcia and Mr. Lynch and I had an excellent work session at MIT's labs on their research in AI and financial services as well.

In all these efforts, we saw demonstrated how AI has the potential to boost efficiency, cut costs, and strengthen the tools used to protect consumers from fraud detection to anti-money laundering but as with any innovation, there are risks. AI systems have to be trustworthy, fair, and secure.

Today we will start that exploration in this Congress on how AI, particularly the emerging capabilities of generative AI, can reshape our financial system.

I look forward to the discussion, and I yield back.

Chairman STEIL. Thank you, Mr. Chairman.

Today we welcome an esteemed panel. We welcome the testimony of Dr. David Cox, vice president of AI Models at IBM Research; Dr. Christian Lau, co-founder and president of Dynamo AI; Mr. Matthew Reisman, director of privacy and data collection at the Center for Information Policy Leadership; and Daniel Gorfine, CEO of Gattaca Horizons LLC; as well as Dr. Nicol Turner Lee, senior fellow and director of the Center for Technology Innovation at the Brookings Institution.

We thank each of you for taking your time to be here. Each of you will be recognized for 5 minutes to give an oral presentation of your testimony. Without objection, your written statements will be made part of the record.

Dr. Lau, you are now recognized for 5 minutes for your oral remarks.

STATEMENT OF CHRISTIAN LAU, CO-FOUNDER AND CHIEF PRODUCT OFFICER, DYNAMO AI

Mr. LAU. Chairman Steil, Ranking Member Lynch, members of the subcommittee, and staff, I thank you for inviting me to testify today, and I am honored to participate in discussions focused on advancing AI across the financial services ecosystem.

In 2021, I founded Dynamo AI alongside my co-founder and CEO, Vaikkunth Mugunthan, during our Ph.D.s at MIT. We started off as a small group of researchers deeply interested in AI but also keenly aware that AI risks pose fundamental challenges to real world adoption.

Since starting as a small group of Ph.D.s, we quickly found ourselves working hand in hand with some of the largest financial in-

stitutions, the most cutting-edge financial technology (fintech) companies, as well as regional banks across America to help them navigate compliance and governance challenges posed by AI.

Today, we not only provide the security and governance layer for many of the largest deployments of AI in banking, but we are also now proudly backed by 40 of the top 100 U.S. financial institutions and a consortium of community banks across America.

Our mission has always been to help enterprises navigate complex regulatory environments, particularly where legal and compliance requirements pose open technology challenges that institutions struggle to solve. We found that, when faced with new technology regulations or internal compliance requirements around new technologies, enterprises are often left paralyzed, asking themselves not only how do I comply with these new requirements but is it even technically possible for us to comply with these new requirements. Nowhere do we see this to be more prevalent than with the struggles of financial institutions striving to adopt AI.

Every day, our team has the ability and opportunity to witness exciting new AI proof of concepts that bring new efficiencies to the bank but for every new exciting AI proof of concept that we encounter, we also see another AI proof of concept fail to make it into production and deliver meaningful value. We commonly see that these projects fail not because the underlying AI technology cannot deliver but, rather, because financial institutions struggle to answer open questions about managing AI risk in heavily regulated, high-impact, and consequential environments.

Some common and, quite frankly, sensible questions about AI risk include: What new security and data leakage risks does AI bring to my organization? As we hand over more autonomy to AI agents, how can we enable those agents to comply with established banking protocol and procedure when carrying out tasks and what happens when an AI assistant inevitably hallucinates or fabricates facts in its response?

While financial institutions often struggle to answer these questions, these are not intractable problems. At Dynamo, we work with a multitude of financial institutions to establish effective AI risk management that accelerates rather than blocks their AI transformation.

To truly manage these new risks and unleash AI innovation, financial institutions and regulators must embrace technology solutions that can help risk and compliance teams scale their oversight, including controls like AI guardrails, red-teaming evaluations, and auditability over AI usage.

Importantly, the AI landscape is evolving at breakneck speed, and regulators and policymakers need to adopt governing frameworks that keep up with the pace of innovation, rather than falling behind on new opportunities and risks.

Regulators should expand the use of AI sandboxes, as called for in the administration's AI Action Plan and the bipartisan H.R. 4801, to not only enable technology teams to experiment with AI on high-impact use cases but also bring leading evaluation and red-teaming technology to rigorously test these experimental AI against real world risks.

Drawing on global examples, such as Singapore's successful sandbox programs, the U.S. can strengthen competitiveness while promoting secure and compliant AI.

Across financial institutions and governing bodies that we speak to every day, we are starting to see comprehensive AI risk management take shape. We believe that this will be key to ushering a new and exciting era of advancement for the industry.

On behalf of the entire team at Dynamo AI, thank you for the opportunity to testify, and I welcome the opportunity to work with the subcommittee to create a competitive, secure, and compliant AI ecosystem within financial services.

Thank you.

[The prepared statement of Mr. Lau follows:]

Written Testimony of Dr. Christian Lau
Co-Founder and President, Dynamo AI

**Digital Assets, Financial Technology, and Artificial Intelligence Subcommittee of
U.S. House of Representatives Financial Services Committee**

**Hearing on “Unlocking the Next Generation of AI in the U.S. Financial System for
Consumers, Businesses, and Competitiveness”**

September 18, 2025

Chairman Steil, Ranking Member Lynch, members of the Subcommittee, and staff, I thank you for inviting me to testify today and am honored to participate in discussions focused on advancing AI across the financial services ecosystem.

Introduction

We are at a pivotal moment in the history and development of technology. Artificial intelligence (AI) stands to reshape so much across our economy and daily lives, and this is particularly true within the financial services sector. With a vibrant and active market of players driving much of this change, it is an honor to be able to contribute to this conversation on behalf of Dynamo AI and the thriving ecosystem of AI and security start-ups in America.

In 2021, I founded Dynamo AI alongside my co-founder and CEO Vaikkunth Mugunthan during our PhDs at MIT. Our mission has always been to help enterprises navigate complex regulatory environments, particularly where compliance requirements pose open technology challenges that institutions struggle to solve. We found that, when faced with new technology regulations or internal compliance requirements around new technologies, enterprises are often left paralyzed, asking themselves not only “how do I comply with this new requirement?” but also “is it even *technically possible* for us to comply with this new requirement?”

Nowhere did we see this to be more prevalent than with the struggles of financial institutions striving to adopt AI. Since founding Dynamo AI, we’ve had the opportunity to work with some of the largest global financial institutions, the most cutting edge fintech companies, as well as regional banks across America to help them navigate compliance and governance challenges posed by AI. Dynamo AI itself is backed by over 40 of the top 100 US financial institutions and a consortium of community banks who often lean on Dynamo AI to navigate their governance of AI and securely deliver AI applications into production.

Every day, our team witnesses exciting new AI proof of concepts (POCs) of financial services providers that promise to transform the customer experience, enable more efficient and comprehensive compliance, or enable more data driven decisioning. But for nearly every exciting AI POC we encounter, we also see another AI POC fail to make it into production and

deliver value. We commonly see these POCs fail not because the underlying AI technology can't deliver, but rather because financial institutions struggle to answer open questions about managing AI risk in heavily regulated environments. We routinely encounter technology teams at financial institutions paralyzed by questions surrounding AI risks from legal, compliance, and security stakeholders: How can institutions teams manage risks of giving every banking associate an AI assistant that can be prompted in infinitely different ways? How can such AI assistants encourage rather than mislead employees in following bank protocols and procedures? What happens when an AI assistant inevitably hallucinates, fabricating facts in its response?

While financial institutions often struggle to respond to these concerns, these are not intractable problems. At Dynamo, we've worked with a multitude of financial institutions to establish effective AI risk management that accelerates, rather than blocks AI transformation. This often first involves institution-wide education about AI, including AI's capabilities, function, and risks, followed by an initiative to align diverse risk stakeholders across the financial institution around a cross-functional governance framework. To layer in the necessary technical controls, financial institutions must also embrace technology solutions that can help risk teams scale.

Our Dynamo team spends just as much time educating legal, risk, compliance, security, technology teams, as well as regulators about best practices in AI governance as we spend on implementing technical controls around AI risks, including controls like AI guardrails, red-teaming evaluations, and observability over AI usage. For example, to date, our AI guardrails service checks over 1 million user interactions every day for security vulnerabilities and noncompliance with bank policies, giving nontechnical risk stakeholder auditability into AI applications across the bank. We're starting to see comprehensive AI risk management take shape, which we believe will be key to ushering in this exciting new era of advancement for financial services.

Advancement of AI across Financial Services

A Proliferation of Economic Advancement and Business Opportunity

Working across the financial services sector, our team has gained unique insight into AI opportunities through two key experiences: leading an AI company backed by a consortium of financial institutions—from credit unions to regional banks and systematically important financial institutions—and helping financial services organizations enable AI use cases within this highly regulated sector.

To date, financial institutions have looked to AI primarily to enhance productivity to derive return on investments, increase worker and operational efficiency, and learn how to establish effective AI governance foundations. This takes many forms, including AI-powered chatbots that help employees understand company policies and business line standards to better execute processes and engage with customers and colleagues. This also is evident in the rapid development of AI used to generate code for technology systems. I expect AI to further extend across the financial services value chain, including financial product operations, investment

analysis, and customer experience over the near term. Additionally, financial regulators themselves are starting to look to AI as a tool to support market oversight and monitoring of potential financial misconduct in the marketplace.

Many institutions have identified hundreds of AI use cases in their internal queues, ready for exploration, assessment, and deployment. This is a result of the ingenuity and excitement of financial services personnel from all lines of business and operations, looking to deploy AI to strengthen their organizations and propel customer value.

The financial services vendor ecosystem is also rapidly introducing AI to enhance core technology and operational processes, from compliance management to procurement. This includes the introduction of AI capabilities within existing third party applications, such as a sales or document processing software that an institution may already use. While this opens up opportunity and efficiencies in existing embedded processes, it also introduces additional institutional risks that we will discuss further.

As this subcommittee learns about AI's benefits and risks in financial services, it is crucial to recognize how community and regional banks can leverage AI to better serve consumers and strengthen local economies. Some of the most forward-thinking community and regional banks I interact with view AI as a transformational, once in a generation opportunity to deepen customer relationships, enhance community engagement, prevent fraud, and achieve operational efficiencies that directly benefit consumers.

AI's ability to personalize services at scale enables community and regional banks to compete more effectively with larger institutions, while delivering tailored financial solutions that expand consumer access and financial wellbeing. This technology can help smaller banks offer sophisticated services previously available only at major institutions, creating a more competitive and consumer-friendly banking landscape.

With proper risk mitigation, AI deployment in financial services can deliver significant advantages for consumers across all institution sizes—from enhanced fraud protection and more scalable customer service to broader access to credit and banking services in underserved communities.

Landscape of AI Risks for Financial Services Institutions and Regulators

I believe the financial services industry is in the midst of a pivotal moment in history with respect to risk management. Alongside the development and deployment of initial AI use cases across the financial services landscape, organizations are learning how to operationally govern AI use and integrate these activities into existing risk management frameworks. I see this in the cross-functional AI governance working groups that are being established, the policies being written and deployed, the training that is expanding across employees, and the active engagement with regulators, trade associations, and the marketplace. More broadly, financial services risk management practices often become the standard that other industries adopt. Therefore, it is critical that policymakers and regulators clearly signal their priorities—both the innovation they

want to encourage and the risks they consider most important to address. This regulatory tone will shape how financial institutions design their AI governance frameworks today and how these systems operate for years to come.

There are risks that are unique to AI, as well as risks that are unique to AI and financial services. For the subcommittee, I will provide a brief overview of a number of these AI risks, of which I can answer questions about, but will also provide deeper insights into five risks that I believe are critical to mitigate to advance a vibrant, secure financial services ecosystem.

Key Cross-Sector AI Risks

Generative AI introduces several unique risks that organizations often struggle to mitigate and monitor, which in turn can delay effective use and integration of the technology. These risks include:

- The risks of **hallucination**, a scenario in which AI systems generate false, misleading, or fabricated information, while presenting it as factual or accurate. This risk can degrade trust and lead to harmful or misinformed decision-making by whoever makes use of AI outputs.
- The risk of **adversarial security attacks**, including prompt injections that manipulate AI responses through malicious inputs and jailbreaking techniques that bypass AI guardrails, which can exploit model vulnerabilities and circumvent defenses at scale, require AI-specific threat evaluations and controls.
- **Data and Access Risks** including vulnerabilities where third party AI providers may consume volumes of enterprise data fed into models, resulting in the potential leakage of sensitive enterprise or client data to third parties.
- The risk of **misuse of AI systems** for unauthorized, high-impact use-cases, such as writing legal documents, obtaining unwarranted investment or healthcare advice, or making creditworthiness assessments without proper checks.

Key AI Risks Impeding Financial Services Innovation

The financial services industry is home to a suite of rigorous risk management practices that aim to protect consumers, introduce effective challenge, independent thought, and advance financial opportunity for our citizens and the marketplace. The worthy challenge of introducing AI in the financial services ecosystem is to balance risk management expectations, while also promoting innovation for consumers, vendors, and our marketplace. Therefore, there are a few key risks I believe require attention from the subcommittee.

One risk is rooted in the **explainability** of a model's response. Simply put, it is impossible to determine exactly how or why a Generative AI model responded the way it did. This, in itself, makes Generative AI models different from conventional models used in financial services today, such as statistical or probabilistic models, for which a decision can be traced back to a specific component of the model. The reality of explainability and lack of traceability in AI

models may warrant a revision of existing risk management guidance, evaluation methods, and subsequent controls.

An additional noteworthy risk is that of protecting **model sovereignty**, which refers to an organization's control over the AI models they deploy, including the data, infrastructure, and capabilities needed to maintain independence from external providers and ensure regulatory compliance. This is of particular importance for financial institutions and regulators to align on, as models are often developed and managed by third parties that may service other institutions and/or critical technology components of the same institution. Importantly, this risk is only amplified if US financial institutions utilize models that are developed or trained in nations with adversarial relationships with the United States, including China. A discussion focused on acceptable third-party risk management protocols is a critical steppingstone to effective internal and external oversight.

A strength of the financial services regulatory system is that institutions can interpret principles-based rules and turn them into practical compliance processes—even deciding when to accept certain risks. With AI now entering processes once handled by people, the third major risk is **whether institutions can keep AI aligned with their own definitions, policies, and risk tolerance—their “ground truth.”** AI tools often come with assumptions from vendors or global bodies, but it is imperative that financial institutions have the ability and controls to maintain autonomy over their “ground truth.” For example, a regional bank may have a specific definition of “legal advice” and strict rules on when it can or cannot be given. An AI system might not share that definition. Without guardrails that let organizations overlay their own rules on third-party AI models, compliance gaps open and competitiveness suffers.

As alluded to earlier, AI is being actively enabled within a host of existing embedded technologies and processes already in place across financial institutions. For example, many financial institutions have human resource software or document processing software that may have AI components enabled, including document analysis and chatbots. However, market or organizational dependence on a single AI provider or tightly coupled ecosystem increases vulnerability to operational risk or systemic failure. The result is an increase in **third- and fourth-party risks**, in addition to risk surrounding **vendor concentration**. Encouraging diversity in infrastructure and support can improve market resilience, financial stability, and reduce consumer harm.

Finally, it is also worth briefly discussing the newest frontier of AI and AI Risk: **Agentic AI**. In short, AI agents are AI systems that can independently execute complex tasks, make decisions, and take actions on behalf of users or organizations. However, even simple AI agents require organizations to access sensitive information and make impactful decisions. And the value derived from replacing manual work with AI reduces human control and oversight over AI decision-making. As a result, Agentic AI amplifies many of the aforementioned risks. The more powerful the AI agent, the more risk there is in its deployment, as this requires organizations to transfer more decision-making authority from humans to AI systems. Therefore, for AI agents to truly provide value, organizations will need to mobilize novel technologies to rigorously test, sandbox, and embed safeguards into these tools.

Dynamo AI's AI Compliance Strategy team has documented an inventory of key AI risks for the subcommittee, including key considerations and strategic implications for policymakers and industry leaders alike. For more, see Appendix A.

AI as an Enabler for Risk Management and Compliance

Despite the risks presented by AI, one of the most noteworthy benefits of this powerful technology is that **AI also serves as a critical protective function, or control, to mitigate many key AI and security risks and an effective enabler for organizations.**

Dynamo is proud to be a leader in this regard, developing technologies that product and risk management teams more comprehensively evaluate and guardrail both AI-specific threats, as well as compliance and security challenges. Within this context, we have seen novel innovations in the AI trust, security, and risk management ecosystem directly enable organizations to accelerate the deployment of AI applications, while maintaining compliance with complex regulatory requirements.

At Dynamo, our product suite uses AI to provide streamlined tests, evaluations, and real-time protections for AI applications launched in the financial services ecosystem.

AI Used to Power Simulated Security Attacks and Evaluations

With AI systems, the threat landscape of adversarial attacks is both constant and expansive. It is virtually impossible—and economically infeasible—for organizations to staff departments with sufficient personnel to effectively predict and block all possible adversarial threats that may target an AI system. Moreover, new attack methods emerge daily, requiring protection methods to be continuously updated and personnel to conduct ongoing alert review and escalation. In my work at Dynamo, our team has **integrated AI into our test and evaluation suite**, using AI to attack existing models to identify threats proactively for human review and decision-making. In other words, the solution to many key challenges in this space lies in leveraging AI itself as a defensive tool.

Real-time AI Guardrails

A core challenge for the financial services marketplace is how to deploy AI while complying with existing laws and regulations. Dynamo has worked with institutions to develop what are known as AI **“guardrails”** – a specific category of technology we have helped to pioneer that can moderate how a model behaves, what data is permissible to enter or leave an AI system, and to do this in alignment with each institution's policies and definitions. In a highly regulated sector where compliance with regulatory guidance and requirements must be embedded into every process, AI guardrails enable financial institutions to scale their compliance interpretations and best practices and enable high-value AI use cases. I believe this to be one of the primary challenges of AI adoption across the financial services ecosystem; but, once guardrails expand across the ecosystem, organizations can fully realize the value and returns of these technologies, alongside a more sound and secure banking ecosystem.

Innovative Tools for AI Monitoring

A key component of effective AI risk management is to institute comprehensive human-in-the-loop observability to give humans the ability to monitor models, interactions, and demonstrate compliance with internal policy and regulations. Observability equips internal risk management and audit teams, as well as regulators with evidence needed to substantiate a level of operational assurance that can be measured and tested. At Dynamo, we've built a full observability suite to allow organizations to continually observe all internal and customer-facing AI interactions, so they can strengthen controls as AI is in use. While a complex technical solution in its own right, AI observability provides organizations with the necessary reporting and alerts for internal monitoring of powerful systems, and it will prove to be an essential ingredient for effective AI oversight in the sector in the near future.

These AI-powered approaches represent a necessary evolution in risk management. Together, they suggest that sustainable AI adoption in financial services will depend on institutions' ability to implement technology-based controls that can operate at the speed and scale of the systems they oversee.

Key Considerations for Fostering a Competitive, Secure Ecosystem for AI

As the Subcommittee weighs policy and technical considerations for the continued promotion of AI innovation and a vibrant and fair financial services ecosystem, there are a few actions I would like to highlight.

Sandboxes as a Vital Tool for Innovative Oversight and Information-Sharing

Regulators, with the participation of financial services institutions, should continue to establish and utilize AI sandbox environments, a concept referenced in the current administration's recent AI Action Plan. These environments, either established by a financial regulatory and/or market consortia, allow parties on both sides of financial markets to explore AI use cases, risks, as well as acceptable controls to mitigate those risks. They also allow innovative American companies such as ours to participate easily alongside the broader vendor technology landscape and showcase advancements. In turn, sandboxes educate a wide variety of stakeholders across regulatory agencies and leaders within financial services institutions.

We applaud Committee Chairman Hill, Representative Torres, leaders of this Subcommittee, and colleagues in the Senate for introducing H.R. 4801, the Unleashing AI Innovation in Financial Services Act. This Bill strikes a strong balance between supporting innovation, fostering governance within enterprises, and educating regulators on emerging AI use cases and risks, allowing companies to experiment with the emerging technology and conduct necessary tests to assess key risks.

Governments across the globe, including in Singapore and the United Kingdom, have used sandboxes to drive technology forward and inform future policy decisions. We would welcome the opportunity to share best practices and insights from our current work in AI sandboxes.

Growing the AI Evaluation Ecosystem

Continue to support a vibrant ecosystem of AI providers, including within the independent AI evaluation and guardrail market, whether through guidance, regulation, or support for effective challenge risk management methodologies. The administration's AI Action Plan underscores the importance of growing this ecosystem, and it is critical that this community includes not only federal agencies but also companies that are incentivized to accurately identify and protect against existing and emerging AI risks. We have seen similar requirements for independent evaluation in the areas of financial crimes and fraud, and I believe similar expectations may be warranted when applying AI security and compliance controls. As the AI industry continues to grow and new players emerge, it is essential to ensure that the fox is not guarding the henhouse.

Not only are AI evaluations important for driving secure AI adoption in the commercial space. They are also vital as federal agencies look to deploy AI. As mentioned in the AI Action Plan's section calling on agencies to "Accelerate AI Adoption in Government," the General Services Administration (GSA) is tasked with creating an "AI procurement toolbox" for federal agencies to easily acquire AI tools. As part of this agency-ready AI marketplace, we recommend the incorporation of adequate independent test and evaluation technologies into the federal procurement process for high-impact AI use cases, such that federal agencies can easily deploy AI in a compliant and secure way.

Considerations for the Future of Model Risk Management

While mechanisms or applied controls for AI model evaluation may differ from historical model risk management, we are energized by ongoing discussions between private and public sector leaders on acceptable controls for AI use cases in financial services. We ask the Subcommittee to continue to promote an open dialogue between standard-setting institutions, financial regulators, and the financial services industry to arrive at best practices for mitigating risks found in common financial services use cases. This may, in turn, lead to new guidance that may be well received by the industry.

Finally, as the American AI industry continues to grow, we would like to underscore the importance of policy and regulatory initiatives that support the advancement of America's start-up ecosystem. This opportunity is a testament to the Subcommittee's commitment to supporting entrepreneurship, advancing America's technological competitiveness, and giving small businesses a platform. For that, I am very grateful.

Conclusion

On behalf of the entire team at Dynamo AI, thank you again for the opportunity to testify. I welcome the opportunity to work with the Subcommittee to further create a competitive, secure, and compliant AI ecosystem within the financial services industry.

Appendix A: Key Frontier Findings: Current and Emerging AI Risks

I. Foundational and Structural Risks

Risk Category	Key Consideration	Strategic Implication
Supply Chain & Infrastructure	A single weak point can compromise systems across industries. Organizations should diagnose vulnerabilities in their AI supply chain and develop standards to secure it.	Vulnerabilities can be embedded throughout the AI pipeline, exposing the AI pipeline to coordinated attacks.
Vendor Concentration	Market or organizational dependence on a single AI provider or tightly coupled ecosystem increases vulnerability to operational risk or systemic failure.	Encouraging diversity in infrastructure and support can improve resilience, reduce long-term exposure to single points of failure, and market interruption.

II. Governance, Behavior, and Control Risks

Risk Category	Key Consideration	Strategic Implication
Independent Validation & Control	People, process, and technology risks may not be identified or mitigated without independent AI validation and complementary incentives.	When oversight mechanisms are tightly coupled to the AI systems they monitor, risk detection or response is constrained. Independent evaluation and guardrail design can improve transparency and build trust, especially when market incentives facilitate risk identification and mitigation.
Monitoring & Detection	Legacy security tools offer limited visibility into AI model behavior. New attacks can exploit the opacity of models and bypass basic monitoring capabilities.	Investment is needed in AI-native monitoring tools and real-time anomaly detection to ensure meaningful AI visibility and response.
Human Oversight	AI systems often lack human oversight or effective workflows for non-technical stakeholders, limiting oversight into AI systems.	Human-in-the-loop oversight and thoughtful human-integrated workflows are essential for ensuring transparency, accountability, and trust.

Appendix A Continued: Key Frontier Findings: Current and Emerging AI Risks

II. Governance, Behavior, and Control Risks

Risk Category	Key Consideration	Strategic Implication
Adversarial Attacks	Prompt injection and jailbreaking can exploit model vulnerabilities and bypass defenses, at scale.	Traditional security measures may be insufficient. Organizations should consider AI-specific threat modeling, red teaming, and continual model hardening.
System Reliability & Security	Hallucinations or deceptive outputs from AI degrade trust and potentially lead to harmful decisions.	Thorough validation, security audits, and hallucination testing are important, especially in high-stakes or public-facing systems.
Alignment & Control	AI may misinterpret prescribed goals and ignore human intent, behaving dangerously and unexpectedly as it scales.	As AI capabilities advance, Governance frameworks should evolve, using accountability in system design, testing, and autonomous behavior.

III. Data and Access Risks

Risk Category	Key Consideration	Strategic Implication
Data Security & Privacy	Models can memorize and leak sensitive data; sophisticated attacks can bypass standard privacy safeguards.	Investing in proactive data governance is necessary to meet evolving regulatory expectations and to create adaptive protections.
Authentication & Identity	AI-generated media can bypass identity checks and can be used to maliciously and convincingly target individuals down the line.	To counter the evolving risks of synthetic media and AI-driven fraud, organizations should invest in improved identity verification systems.

Chairman STEIL. Thank you very much.
 Dr. David Cox, you are now recognized for 5 minutes.

**STATEMENT OF DAVID COX, VICE PRESIDENT, AI MODELS;
 IBM DIRECTOR, MIT-IBM WATSON AI LAB**

Mr. COX. Chairman Steil, Ranking Member Lynch, and distinguished members of this committee, thank you for the opportunity to testify. My name is David Cox, and I serve as the vice president for AI Models at IBM Research and the IBM director of the MIT-IBM Watson AI Lab.

While the excitement around AI has intensified in recent years, artificial intelligence has fascinated researchers for many decades. To give you a sense of our historical place in this, IBM has been on the vanguard of this ongoing revolution from the very beginning, when an IBM researcher co-authored the proposal for the 1956 workshop, Dartmouth conference, that coined the term “artificial intelligence” and gave the field its name.

IBM also has a long history of helping enterprises, including in the financial sector, use artificial intelligence technologies to unlock business value, and doing so in ways that are responsible and engender trust.

Our work spans a broad spectrum, from helping to identify appropriate use cases, to providing tooling to govern both the development and deployment of AI systems, to inventing new technologies for trustworthy AI.

From the start, IBM has tested AI internally before offering it to others, an approach we call Client Zero. By deploying AI internally, we ensure our models are stress tested in the real world environments before they ever reach clients.

For example, by leveraging IBM’s watsonx.ai and automation tools, we help to augment the skills of our own workforce by eliminating repetitive tasks. This has enabled our employees to focus on more challenging, rewarding, and impactful work with the time they save.

For the financial industry, the implications of AI, particularly generative AI and large language models, are transformative. Learning Management System (LMS) are not merely chatbots, they are multifunctional tools that can be adapted across a wide range of financial services use cases. These opportunities come with challenges. Large models consume enormous computing resources, raising both costs and concerns about energy use.

Transparency is vital. Enterprises and regulators must understand the provenance and quality of the data underlying deployed systems. Security is also paramount as organizations seek to safeguard sensitive information when using cloud-based systems but perhaps the greatest risk is hesitation. If industry delays too long, consumers and the U.S. economy may miss out on the early benefits of adoption.

Responsible governance is not a break on innovation. It is a mechanism that ensures that innovation can be deployed securely and sustainably.

In regulated environments, firms must understand exactly what data underpins their models and be able to audit those systems

over time. That is why IBM's enterprise AI efforts are grounded in three principles: open, trusted, and secure.

Open-source AI strengthens transparency, produces dependence on proprietary vendors, and enhances U.S. competitiveness. Trust is built through transparency and data curation, training processes, and lineage between models and data. Finally, security must be embedded throughout the AI life cycle, from data collection to deployment, backed by continuous oversight rather than reactive compliance.

With this in mind, I urge policymakers to support open ecosystems, to regulate applications rather than technologies in the abstract, and to promote transparency requirements that allow enterprises and regulators alike to test models, evaluate accuracy, and understand safeguards. These steps will foster an environment where innovation can thrive responsibly.

AI is not about replacing people but augmenting them. It is about empowering professionals, enriching consumers, and expanding opportunity. By embracing openness, insisting on transparency, and embedding security, we can ensure that AI strengthens both the financial system and America's global competitiveness.

Thank you, and I look forward to your questions.

[The prepared statement of Mr. Cox follows:]

Testimony of Dr. David Cox
Before the HOUSE FINANCIAL SERVICES COMMITTEE SUBCOMMITTEE ON DIGITAL
ASSETS, FINANCIAL TECHNOLOGY, and ARTIFICIAL INTELLIGENCE

Hearing on “Unlocking the Next Generation of AI in the U.S. Financial System for
Consumers, Businesses, and Competitiveness”
Thursday, September 18th, 2025

Chairman Steil, Ranking Member Lynch, and distinguished members of the Committee, thank you for the opportunity to appear before you today. My name is David Cox, and I serve as Vice President for Artificial Intelligence Models at IBM Research and Director of the MIT-IBM Watson AI Lab.

A Brief History of AI: From Rules to Watson

For decades, artificial intelligence has captured the human imagination. Ever since the term was first coined by computer scientist John McCarthy in a 1955 proposal for the 1956 Dartmouth Summer Research Project on Artificial Intelligence, IBM has played an integral role in this technology’s evolution. Indeed, it was an IBM researcher, Nathaniel Rochester, who co-authored the Dartmouth proposal and co-hosted the historic conference. But it wasn’t until the 1970s, however, that much of the technology’s practical evolution began unfolding.

Early AI systems were largely rules-based “expert systems.” These programs relied on manually encoded knowledge, painstakingly crafted by engineers, to make decisions. While useful in narrow contexts — like medical diagnosis support or industrial automation — these systems were brittle, unable to adapt beyond their programmed rules.

By the 1980s and 1990s, machine learning emerged. Instead of encoding all the rules by hand, researchers trained systems using statistical methods and data. This shift gave rise to algorithms for fraud detection, credit scoring, and basic speech recognition — tools that began to find homes in finance, healthcare, and telecommunications. It was during this time, in 1998, that IBM’s “Deep Blue” made headlines by defeating Garry Kasparov, the world chess champion. But chess was only a milestone. It demonstrated how computing power, when combined with intelligent design, could solve problems long thought unsolvable.

The 2000s ushered in the era of scale. With exponential increases in computational power and data availability, machine learning accelerated dramatically. In 2011, IBM Watson became a household name by winning *Jeopardy!*, defeating two of the game’s greatest champions. This was not just a media spectacle; it was proof that computers could understand natural language, navigate ambiguity, and reason across vast stores of information. Watson symbolized the leap from AI as an experimental tool to AI as a practical enterprise technology.

Watson’s success laid the foundation for IBM’s modern AI efforts. What began as a demonstration on a television stage evolved into a full suite of enterprise-ready tools, now

unified under watsonx®, which combines foundation models, generative AI, and trusted governance.

IBM's Watson Legacy and the "Client Zero" Approach

From the very beginning, IBM has treated Watson not as a product in isolation, but as a living ecosystem. We have applied Watson technologies inside IBM's own operations, making ourselves the first and most demanding client — what we call "Client Zero."

AI tools, models, and governance frameworks are tested in IBM's own systems before reaching clients. This philosophy means that when we speak to our clients in banking, insurance, and capital markets about AI, we speak from firsthand experience, not just theoretical rumination.

For example, in 2024 alone, AI-powered automation tools deployed internally saved IBM employees an estimated 3.9 million hours of routine work. That time was reinvested into higher-value activities, from research and product development to client services. By using AI-powered financial transparency tools, IBM also identified cost-saving opportunities that contributed to nearly \$600 million in IT savings since 2022. By 2025, our Client Zero initiatives are on track to generate approximately \$4.5 billion in productivity improvements, spread across more than 70 business functions, including finance.

This internal proving ground allows us to deliver AI solutions to clients with confidence: if they can scale inside IBM, they can scale anywhere.

Partnerships and Open Innovation

IBM has also recognized that no one company can solve AI's challenges alone. That is why our AI story is also a story of partnerships, with academia, others in industry, and the larger open source community.

Academia. Through the **MIT-IBM Watson AI Lab**, of which I am the director, we collaborate with leading researchers to advance fundamental AI science. This lab is one of the largest university-industry AI research collaborations in the world, and it has produced breakthroughs in natural language processing, trustworthy AI, and domain-specific model development.

Industry. We partner with financial institutions to co-develop AI models optimized for compliance, customer engagement, and fraud detection. By embedding client expertise into model training, we ensure that Watson is not a "one-size-fits-all" solution, but a toolkit tailored to highly regulated, high-stakes environments.

Open Source. IBM has been a champion of open innovation. Just as Linux became a backbone of modern computing, open source AI models can become a foundation for trustworthy, transparent, and secure enterprise AI. By contributing to open source communities, IBM ensures that regulators, academics, and startups have the same opportunity to inspect, adapt, and improve upon foundational models.

These partnerships reflect IBM's longstanding ethos: that AI should be a collaborative, democratizing force, not a black box controlled by a small handful of providers.

Financial Services: Watson in Action

Nowhere is IBM's AI impact clearer than in financial services — an industry that thrives on data, trust, and efficiency. IBM and our partners are applying Watson-powered AI solutions across a wide array of domains and applications.

Regulatory Compliance and Auditability. In the regulatory compliance and auditability realm, AI can automatically generate and maintain auditable trails of financial activities, helping compliance officers navigate complex and evolving regulatory requirements. Long prospectuses and contracts can be summarized in minutes, giving decision-makers clarity without sacrificing detail.

Customer Service and Personalization. Financial firms are deploying IBM-enabled chatbots and natural language systems to deliver faster, more personalized customer service. Clients can interact with their institutions using non-technical language, resolving routine queries instantly while freeing human representatives for more complex cases.

Fraud Detection and Risk Management. AI helps financial institutions detect anomalies in real time, flagging fraudulent activity before it spreads. By analyzing vast amounts of transaction data, domain-specific models can make credit, investment, and risk-based decisions in microseconds.

Investment Research. AI tools are accelerating investment research by scanning global news, financial statements, and regulatory filings, producing concise and actionable insights. This not only saves time but levels the playing field by giving smaller firms access to capabilities once reserved for the largest institutions.

IT Resilience and Developer Productivity. Financial institutions use watsonx to optimize IT operations, identify cost inefficiencies, and boost developer productivity. This leads to more resilient infrastructure, faster innovation cycles, and improved consumer trust in the reliability of financial systems.

Each of these applications builds on IBM's central promise: to augment human expertise, not replace it. By working alongside financial professionals, AI enhances decision-making, accelerates workflows, and deepens trust with customers.

Risks and Challenges

As with any transformative technology, there are challenges we must anticipate and manage even as we take advantage of AI's opportunities. Many of the core principles remain unchanged, such as the importance of evaluating risks in relation to specific use cases rather

than the technology alone. Yet large language models (LLMs) introduce new dimensions that demand careful attention.

One area of concern is scale. LLMs have been growing rapidly in size, and larger models require significantly more computational resources to train and run. This drives up both costs and power consumption. Energy demand in particular could put real pressure on infrastructure. Some projections suggest that by 2040, the global computing sector's energy use may surpass the world's entire energy budget.

Another pressing issue is transparency. Training LLMs requires massive amounts of data, but opacity can create serious challenges for enterprises, especially those operating under strict regulatory oversight, where both companies and regulators need confidence in the provenance and quality of the data underlying deployed systems.

Security is also top of mind. Organizations worry about safeguarding proprietary and customer data, particularly when using LLMs delivered through cloud-based services. Some firms have prohibited employees from using certain as-a-service models out of concern that sensitive information could be exposed to third-party vendors without sufficient security guarantees.

Perhaps the most consequential risk, however, lies not in overuse but in underuse. If industry hesitates too long to deploy AI, the broader economy and consumers may miss out on the early advantages of adoption. The path forward requires accelerating adoption while ensuring strong governance frameworks. Responsible AI is not a brake on innovation, but a catalyst for realizing its benefits securely and sustainably.

Responsible AI: Open, Trusted, Secure

The financial sector understands better than most that trust is painstaking to earn yet can vanish in an instant. In highly regulated industries, firms should have confidence in their models and maintain the ability to verify those systems over time. This is why IBM's enterprise AI strategy rests on three core principles: **Open, Trusted, and Secure.**

Open. We are firm believers in the importance of open source AI. Just as Linux became a backbone of the modern digital economy, open technologies drive security, collaboration, and innovation. Open models give regulators, researchers, and businesses the ability to examine, refine, and tailor AI to their particular needs. They also reduce dependency on a small group of closed, proprietary vendors — a key factor for safeguarding national security, competitiveness, and long-term resilience.

Open source AI is not a liability but one of the strongest tools we have to manage risk. It advances transparency, auditability, security, and cost-efficiency. IBM contributes actively to the open source large language model ecosystem and supports open data governance practices with this philosophy in mind.

Moreover, open innovation strengthens the entire AI landscape. It promotes economic dynamism, bolsters security, and aligns with democratic values by making AI development more transparent and collaborative. It enables smaller firms and academic institutions to participate without prohibitive upfront costs and helps cultivate the workforce America will need to stay ahead in the global AI race. By pooling collective expertise through open contribution, we can ensure the AI future benefits everyone — not just a select few.

Trusted. Transparency is the foundation of trust. Preparing the data that powers enterprise-grade language models is a serious responsibility. IBM applies AI-based content filters, extensive blocklists, and curated datasets to reduce the likelihood of harmful material contaminating our models. We also deliberately add authoritative sources to ensure our models are built on high-quality, reliable data.

We make the training of our Granite models transparent, disclosing the data processes involved and releasing the frameworks we use to curate inputs. Our custom data management system maintains full lineage between models and their training data. This system ties directly into IBM's internal clearance process. Stanford University's annual transparency index ranks IBM's models among the most open of any peer provider.

Users of AI systems — especially those handling sensitive information — deserve to know what is inside the models they rely on.

Secure. Security must be embedded throughout the AI lifecycle — from data gathering and preprocessing, to model training and development, to deployment and usage. A robust, risk-based approach secures not only the data and the models, but also the infrastructure they depend upon. This includes access controls, encryption, anomaly detection, and machine learning detection and response (MLDR) capabilities designed to defend against evolving threats.

IBM's integrated governance program (IGP) puts this into practice. It shifts away from reactive compliance toward continuous oversight across data, privacy, security, and AI systems. By embedding compliance into day-to-day operations, IBM is able to deliver AI solutions faster while maintaining the highest levels of trust and accountability. The IGP itself operates on IBM's own technology, serving as a proving ground for the secure, trusted AI practices we bring to clients.

Policy Recommendations

To foster an environment that balances innovation with appropriate safeguards, policymakers should prioritize policies that emphasize openness and technological progress. A healthy, competitive AI ecosystem depends on open source AI to ensure America's future in this domain is shaped by the many, not controlled by a select few. As the recent AI Action Plan noted:

“We need to ensure America has leading open models founded on American values. Opensource and open-weight models could become global standards in some areas of business and in academic research worldwide. For that reason, they also have geostrategic value. While the decision of whether and how to release an open or closed model is fundamentally up to the developer, the Federal government should create a supportive environment for open models.”

At the same time, policies should avoid unnecessary or overly burdensome regulation. Instead, regulators should leverage existing authorities to address specific AI use cases. Where guardrails are warranted, they should be risk-based and tailored to the distinct roles and responsibilities of organizations across the AI development lifecycle.

To help AI thrive responsibly in financial services, we recommend that policymakers:

- **Support Open Ecosystems.** Avoid unduly burdening the open source AI ecosystem with legislative and regulatory policy proposals that hamper contributions to, and the flourishing of, open innovation. Open source AI models should not be disadvantaged relative to proprietary AI models
- **Adopt Use-Case-Based Regulation.** Regulate applications, not just technology. An AI model used for summarizing SEC filings should be governed differently than one used for autonomous trading.
- **Promote Transparency Requirements.** Enterprises and regulators must be able to answer: Can I test my model for accuracy? What decisions did it influence? What safeguards are in place?

Conclusion

Generative AI is no longer a future technology — it is here, and it is reshaping financial services today. IBM’s journey with Watson, from expert systems to *Jeopardy!* to watsonx, illustrates that AI’s future is not about machines replacing people. It is about machines augmenting people — empowering professionals, enriching consumers, and expanding opportunity.

The financial industry has a once-in-a-generation opportunity to lead by example: to show that AI can deliver value while upholding trust. By embracing openness, demanding transparency, and insisting on security, we can ensure that AI strengthens, not undermines, the financial system.

Thank you, and I look forward to your questions.

Chairman STEIL. Thank you very much.
Mr. Reisman, you are now recognized for 5 minutes.

**STATEMENT OF MATTHEW REISMAN, DIRECTOR, PRIVACY
AND DATA POLICY, CENTER FOR INFORMATION POLICY
LEADERSHIP**

Mr. REISMAN. Thank you, Chairman Steil, Ranking Member Lynch, and members of the subcommittee, for the opportunity to speak with you today.

I am Matthew Reisman, director of Privacy and Data Policy at the Center for Information Policy Leadership, or CIPL, a data and privacy policy think tank within the Hunton law firm, whose mission is to advance best practices for the responsible and beneficial use of data. CIPL facilitates constructive engagement between business leaders, data governance experts, regulators, and policy-makers around the world.

AI plays a critical role in the financial services industry, as this committee documented well in the staff report of its bipartisan Working Group on Artificial Intelligence. For years, machine learning has strengthened financial services institutions' ability to combat fraud, provide richer and more tailored services to existing customers, and extend services to new ones. More recently, generative AI has boosted productivity across functions, from software development to customer service and we are now in the early days of agentic AI, which shows promise for enhancing the experiences of businesses and customers alike. Potential applications are extensive, from streamlining Know Your Customer processes, to back-office operations like payroll and invoicing, to online banking and agentic commerce.

AI's use in financial services also carries risks to consumers, institutions, and the financial system. The bipartisan staff report documents these risks well. Agentic AI may accentuate some risks while at the same time enhancing risk management capabilities in areas such as privacy and cybersecurity.

To secure the advantages of AI within the financial system, we have three broad sets of recommendations:

First, with respect to regulation, pursue a risk-based approach that focuses on the outcomes to be achieved. Avoid overly prescriptive measures, build upon existing foundations, including regulations, guidance, and standards that are already in place. When necessary, clarify or adapt their application to emerging technologies. Consider risks and benefits in equal measure. Incentivize organizations to adopt accountable practices. Build trust through meaningful transparency. CIPL has long underscored that these concepts of organizational accountability are central to smart governance of data and technology.

Second, enable the responsible use of data for model training and development. To function safely, effectively, and fairly, models must be trained and tested using rich datasets. Regulators should apply data protection principles in ways that ensure the quality of AI systems while preserving individuals' privacy. Privacy-enhancing technologies, or PETs, can reduce risks associated with the use of personal data. Examples include synthetic data, which is artificial data that mimics the value of real world data, and differential

privacy, where random noise is added to datasets to prevent identification of any individual's data. Policymakers should encourage continued research on and use of PETs.

Third, engage in cooperative dialog among regulators, technologists, and industry. As AI continues to evolve rapidly, stakeholders must learn from each other. Fostering such exchanges is the heart of CIPL's mission, and regulatory sandboxes offer an invaluable avenue for such dialog. Numerous jurisdictions have established regulatory sandboxes over the past decade, from the U.K. to Singapore to U.S. States, like North Carolina and Delaware. We commend steps to promote sandboxes under America's AI Action Plan, the proposed bipartisan, bicameral Unleashing IA Innovation and Financial Services Act, and other proposed legislation.

CIPL has published numerous papers on the aforementioned topics and is continuing our research on them.

We look forward to the discussion today and to supporting your efforts to secure the benefits of AI and financial services for everyone. Thank you.

[The prepared statement of Mr. Reisman follows:]

Testimony of Matthew Reisman
Centre for Information Policy Leadership

September 18, 2025

Thank you, Chairman Steil, Vice Chairman Emmer, Ranking Member Lynch, and members of the Subcommittee for the opportunity to be with you today. I am Matthew Reisman, Director of Privacy and Data Policy at the Centre for Information Policy Leadership, or CIPL, a data and privacy policy think tank within the Hunton law firm whose mission is to advance best practices for the responsible and beneficial use of data in technology and society. CIPL facilitates constructive engagement between business leaders, data governance experts, regulators, and policymakers around the world.

AI plays a critical role in the financial services industry, as this Committee documented well in the staff report of its bipartisan Working Group on Artificial Intelligence. For years, machine learning has strengthened financial services institutions' ability to combat fraud, provide richer and more tailored services to existing customers, and extend services to new ones. More recently, generative AI has boosted productivity across functions, from software development to customer service. And we are now in the early days of agentic AI, which shows promise for enhancing the experiences of businesses and customers alike. Potential applications in the sector are extensive, from streamlining Know Your Customer processes, to back-office operations like payroll and invoicing, to online banking, to shopping powered by agentic commerce.

AI's use in financial services also carries risks – to consumers, the institutions using the technology, and the financial system. The bipartisan staff report documents these risks well. Like AI solutions before it, agentic AI may accentuate some risks while at the same time enhancing our risk management capabilities, in areas such as privacy and cybersecurity.

To responsibly secure the advantages of AI within the financial system, we have three recommendations:

First, with respect to regulation, pursue a risk-based approach that focuses on the outcomes to be achieved. Avoid overly prescriptive measures. Build upon existing foundations, including regulations, guidance, and standards already in place. When necessary, clarify or adapt their application to emerging technologies. Consider risks and benefits in equal measure. Incentivize organizations to adopt accountable practices. Build trust through meaningful transparency. CIPL has long underscored that these concepts of organizational accountability are central to smart governance of data and technology.

Second, enable the responsible use of data for model training and development. To function effectively, safely, and fairly, models must be trained and tested using rich datasets. Regulators should apply data protection principles in ways that ensure the

quality of AI systems, while upholding the privacy of individuals. Privacy-enhancing technologies, or PETs, can reduce risks associated with the use of personal data. Examples include synthetic data—artificial data that mimics the values of real-world data—and differential privacy, where random “noise” is added to datasets to prevent identification of any individual’s data. Policymakers should encourage continued research on and use of PETs.

Third, engage in cooperative dialogue among regulators, technologists, and industry. As AI continues to evolve rapidly, stakeholders must learn from each other to harness technology’s benefits and mitigate risks. Fostering such exchanges is the heart of CIPL’s mission, and regulatory sandboxes offer an invaluable avenue for them. Numerous jurisdictions around the world have established regulatory sandboxes for technology over the past decade, from the UK, to Singapore, to U.S. states like North Carolina and Delaware. We commend steps to promote sandboxes under America’s AI Action Plan; the proposed bipartisan, bicameral Unleashing AI Innovation in Financial Services Act; and other proposed legislation.

CIPL has published numerous reports which explore these topics in greater detail. They can be accessed through our website (www.informationpolicycentre.com).

Chairman STEIL. Thank you very much.
Mr. Gorfine, you are now recognized for 5 minutes.

STATEMENT OF DANIEL GORFINE, FOUNDER & CEO, GATTACA HORIZONS; FORMER CHIEF INNOVATION OFFICER & DIRECTOR OF LABCFTC

Mr. GORFINE. Thank you, Chairman Steil, Ranking Member Lynch, and members of the subcommittee, for the opportunity to testify before you today.

I am the founder and CEO of Gattaca Horizons, an adjunct professor at the Georgetown University Law Center, and the former chief innovation officer at the U.S. Commodity Futures Trading Commission (CFTC).

Today's topic on unlocking the next generation of AI in our financial system is critically important. The U.S. holds significant competitive and first mover advantages, and we must foster continued development through thoughtful policy approaches. AI presents tremendous opportunities to further expand access, lower costs, and increase efficiencies and competitiveness while also enhancing compliance and regulatory oversight.

To level set, it is important to recognize that AI and financial services are not new. It is part of a steady progression of automation that began decades ago. Today, recent advances in generative AI and agentic AI are creating new possibilities, ranging from generating code and new content to planning and executing complex tasks.

Responsibly developed AI is, indeed, already yielding tremendous benefits. AI can provide more accurate and efficient decision-making, detect patterns that traditional approaches would miss, help regulators keep pace with digital markets, and promote financial inclusion by unlocking credit for historically underserved populations. It is further enhancing customer service, helping to identify patterns of financial crime and market manipulation, and making financial advice more accessible and lower cost for Americans.

As with any area of innovation, however, there are risks associated with AI, including the potential for perpetuating bias, relying on poor quality data, failing to operate as expected, and advancing fraud and scams. The mere speculative potential or fear of future harm, however, should not broadly block development of AI, including by those small firms and community banks seeking to remain competitive in an increasingly digital economy.

To this end, as a key guiding principle, I would encourage everyone to assess new AI-based models on their ability to improve off of a highly imperfect status quo. This principle should apply across AI applications, since a singular focus on risk can blind us to the greater benefits as compared to legacy approaches.

With respect to existing policy frameworks, the financial services industry is well-equipped to manage new technologies and should be a model for other sectors. For decades, a robust technology-neutral regulatory framework has governed the adoption of emerging technologies.

For example, consumer protection laws bar discrimination in lending, whether the decision is made by a human or by an algorithm. In our capital markets, rules against fraud and manipula-

tion, along with circuit breaker technologies, help to mitigate risks related to potential AI-driven trading activity and our financial regulators have long provided robust risk management guidance, principles, and frameworks that address model IT third-party and enterprise risks.

The overarching framework to ensure safe and responsible adoption of AI is in place. The challenge now is thoughtful application through sound, informed, and consistent regulation.

To ensure the U.S. maintains its leadership and competitiveness, I accordingly offer five recommendations:

First, regulators must support, not block, responsible AI adoption. We often hear that innovation is blocked by regulators at the examination and supervisory levels due to vague expectations and endless inquiries that lack clear paths to compliance. Examiners should be well-versed in the benefits and risks of AI and provide the marketplace with clear expectations.

Second, we must ensure financial regulators have the in-house expertise and tools to keep pace with technology. A recent Government Accountability Office (GAO) report found a significant lack of technology skills at Federal financial regulators, which is why codifying innovation offices and equipping regulators with their own AI expertise, tools, and capabilities are essential for effective oversight.

Third, Congress should establish a Federal data privacy framework and ensure open access to quality permissioned financial data. The quality of AI model outputs is inherently tied to the quality of data inputs. Congress should work to establish a national framework that governs data privacy and ensures that consumers have control over how their data is being shared and used.

Fourth, Congress should prevent State laws from undermining Federal financial regulatory frameworks. Given the national nature of AI development and its application in the well-regulated financial services industry, a patchwork of State laws can create conflict, ambiguity, and confusion.

Finally, regulators should further clarify risk management frameworks and promote the value of well-crafted standards. This includes making clear that the mere use of AI, even GenAI or agentic AI, does not inherently make an activity higher risk. Regulators can further help their efforts to keep pace with technological change by supporting well-crafted industry standards through regulatory recognition and safe harbors.

Thank you, and I am happy to answer any questions you may have.

[The prepared statement of Mr. Gorfine follows:]

WRITTEN STATEMENT OF

DANIEL S. GORFINE

CEO, GATTACA HORIZONS LLC; ADJUNCT PROFESSOR OF LAW,
GEORGETOWN UNIVERSITY LAW CENTER; FORMER CHIEF INNOVATION OFFICER
AND DIRECTOR, LABCFTC AT THE U.S. COMMODITY FUTURES TRADING
COMMISSION (CFTC)

BEFORE THE U.S. HOUSE FINANCIAL SERVICES
SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL TECHNOLOGY, AND ARTIFICIAL
INTELLIGENCE

“Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses,
and Competitiveness”

Thursday, September 18, 2025
2:00p.m. EST

Thank you, Chairman Steil, Ranking Member Lynch, and members of the Subcommittee for the opportunity to testify before you today. I am the founder and CEO of Gattaca Horizons LLC, an advisory firm, an adjunct professor at the Georgetown University Law Center, and the former chief innovation officer and director of LabCFTC at the U.S. Commodity Futures Trading Commission (CFTC).¹

The topic of today's discussion is "Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness." This is an important topic and one that has gained renewed and prominent interest given recent technological developments and advances in the field, including with respect to generative AI ("GenAI") and so-called "Agentic AI." It is an area of global innovation and fierce competition where the U.S. holds a number of competitive and first-mover advantages,² and an area that should be responsibly fostered through thoughtful policy approaches. While U.S. markets and financial services have long led the world, there are ample opportunities to further expand access, lower costs, increase efficiencies and competitiveness, enhance compliance and regulatory oversight, and improve customer choice, satisfaction, and opportunity. I accordingly commend the Subcommittee for engaging on this topic, and support deliberative and grounded efforts to understand the longstanding role of AI in financial services, recent technological developments, existing legal and regulatory frameworks, and potential policy approaches that can foster the tremendous opportunity presented, while mitigating risks.

My testimony will track these key topics. I will begin by briefly discussing the evolution of AI in financial services, including a discussion of opportunities and risks, and touch on recent advances. I will then provide an overview of existing legal and regulatory frameworks that have long governed the development and adoption of emerging technologies, including AI, in the financial services sector. And I will conclude by offering policy recommendations that can help to inform this Committee's important work and enhance the capabilities of our financial regulators and competitiveness of our financial system.³

¹ My professional associations are set out in my biography attached as Appendix A.

² See Kahl, C.H and Mitre, J. (2025) "The Real AI Race: America Needs More Than Innovation to Compete With China", *Foreign Affairs*. Available at:

<https://www.foreignaffairs.com/united-states/china-real-artificial-intelligence-race-innovation>; See also Schechner, S., MacMillan, D., and Lin, L. (2018) *U.S. and Chinese Companies Race to Dominate AI*. Available at: <https://www.wsj.com/articles/why-u-s-companies-may-lose-the-ai-race-1516280677>.

³ Content in the first two sections below is derived from [my prior testimony](#) on "AI in Financial Services" before the U.S. Senate Banking Committee in 2023.

I. AI in Financial Services: A Story of Increasing Impact

To level set, AI in financial services is not new,⁴ and instead should be thought of as part of a steady progression of using computers and advanced analytics systems to increase automation in the sector.⁵ Adoption of AI in financial services became clearly identifiable in the 1980s with early applications focusing on investment data analytics,⁶ fraud detection in the 1990s,⁷ and a number of use cases in the 2000s, including around consumer underwriting, risk management, compliance, and cybersecurity.⁸

Over the decades, AI-related technologies have continued to develop, as have potential financial services applications—most recently, we have witnessed the public, open source roll-out of AI models referred to as large language models (LLMs), which are capable of ingesting and learning from large natural language data sets. LLMs are a subset of generative AI models (or “GenAI”), which are capable of creating different forms of new content by publishing probabilistic answers to queries based on prior and ongoing learnings.⁹

Even more recently, public mindshare has become increasingly captivated by the potential of Agentic AI. Agentic AI is built on a foundation of GenAI, but is focused on empowering the AI to proactively plan, make decisions, and execute on such planning and decisions.¹⁰

AI applications in financial services have already yielded tremendous benefits to consumers, small businesses, market participants, and service providers, and regulators tasked with supervising financial institutions, protecting consumers, and ensuring market integrity. By processing large data sets, AI tools can offer predictive insights and analytics that allow for more accurate, efficient, and low-cost decision-making, including in the context of determining

⁴ There is no single, agreed definition of Artificial Intelligence (AI). The OCC defined AI as “the application of computational tools to address tasks traditionally requiring human analysis.” It further noted that “[m]achine learning, a subcategory of artificial intelligence, is a method of designing a sequence of actions to solve a problem that optimizes automatically through experience and with limited or no human intervention.”

See Office of the Comptroller of the Currency (OCC) (2021), *Model Risk Management Version 1.0*. Available at: <https://www.occ.gov/publications-and-resources/publications/comptrollers-handbook/files/model-risk-management/index-model-risk-management.html> (Citing Financial Stability Board (2017) *Artificial Intelligence and Machine Learning in Financial Services: Market Developments and Financial Stability Implications*.)

⁵ Gorfine, D. (2019) *Remarks of Daniel Gorfine, Director of LabCFTC at ISDA, LabCFTC: Developments and Discoveries*. Available at: <https://www.cftc.gov/PressRoom/SpeechesTestimony/opagorfine3>.

⁶ See Stevens, P. (2019) *The secret behind the greatest modern day moneymaker on Wall Street: Remove all emotion*. Available at:

<https://www.cnn.com/2019/11/05/how-jim-simons-founder-of-renaissance-technologies-beats-the-market.html>.

⁷ Christy, C. (1990) *Impact of Artificial Intelligence on Banking*. Available at:

<https://www.latimes.com/archives/la-xpm-1990-01-17-fi-233-story.html>.

⁸ Oliver Wyman (2019) *Artificial Intelligence Applications in Financial Services: Asset Management, Banking and Insurance*. Available at:

<https://www.oliverwyman.com/content/dam/oliver-wyman/v2/publications/2019/dcc/ai-app-in-fs.pdf>.

⁹ Nield, D. (2023) *How ChatGPT and Other LLMs Work—and Where They Could Go Next*. Available at:

<https://www.wired.com/story/how-chatgpt-works-large-language-model/>.

¹⁰ Google Cloud (no date) *What is agentic AI?* Available at: <https://cloud.google.com/discover/what-is-agentic-ai>

creditworthiness when a traditional credit score may preclude access.¹¹ Importantly, especially in higher-risk applications, humans are often “in the loop,” with AI helping to inform, augment, and improve their work or decision-making.

AI is further capable of identifying patterns and anomalies that traditional approaches would be incapable of detecting, including in the context of identifying potential financial crime, fraud, or detecting illegal trading behaviors. Responsibly developing and adopting these tools can help ensure that the U.S. maintains the deepest and best-regulated markets in the world, as well as remains at the forefront of financial services innovation capable of serving our national economic needs and those of American consumers and businesses.

Today, all stakeholders in the financial services space, including banks, fintechs, financial markets firms, and regulators, use different forms of AI in their activities and operations. AI tools are being used to support risk management, compliance and transaction monitoring, trade surveillance, cybersecurity and intrusion detection, trading activity, market intelligence, digital investment advisory, financial crime/AML detection, customer service, and consumer finance underwriting. Increasingly, these same functions are leveraging GenAI and Agentic AI technologies, especially in the context of customer services, fraud detection, cybersecurity monitoring, intelligent document processing, internal coding, and information analysis.¹²

More broadly, consumer underwriting is one of the most prominently discussed areas for AI application given its potential to increase financial inclusion and its risk of perpetuating or creating new forms of bias—though, as noted further below, it is an area that has important consumer protections in place that must be followed regardless of the technology being deployed. It is well known that many legacy scoring systems contain embedded bias and are known to correlate with protected class characteristics.¹³ AI-based underwriting holds substantial promise in—and already is—unlocking more fair and accurate scoring that can expand access to financial services for historically underserved populations. Depending on model design, these models can be more fair and transparent than legacy scoring by quantifying the relative significance of data inputs and assisting the search for less discriminatory models.¹⁴

¹¹ See Oliver Wyman (2019)

¹² Levitt, K. (2025) ‘AI On: How Financial Services Companies Use Agentic AI to Enhance Productivity, Efficiency and Security’, *Nvidia*. Available at:

<https://blogs.nvidia.com/blog/financial-services-agentic-ai/#:~:text=With%20advancements%20in%20agentic%20AI, cost%20savings%20and%20operational%20efficiency>

¹³ Klein, A. (2020) *Reducing bias in AI-based financial services*. Available at:

<https://www.brookings.edu/articles/reducing-bias-in-ai-based-financial-services/>

¹⁴ See generally FinRegLab (2023) *Explainability & Fairness in Machine Learning for Credit Underwriting: Policy & Empirical Findings*. Available at:

https://finreglab.org/wp-content/uploads/2023/07/FinRegLab-Machine-Learning-Research-Policy_Empirical-Overview-FACT-SHEET_July-2023_FINAL-1.pdf

This example is helpful in highlighting a key guiding principle that I would encourage the Subcommittee and regulators to consider when assessing new AI-based models: such models should be judged based on whether they improve off of a highly imperfect—and entrenched—status quo.¹⁵ This principle should apply across AI applications and use cases since a singular focus on risk can blind us to the greater benefits that may be present as compared to legacy approaches.

AI is also benefiting consumers, investors, and market participants with respect to investment activity and financial advice. While the use of algorithms by institutional investors is not new, many new entrants and fintech firms are leveraging AI technologies to make financial advice more accessible, transparent, and lower cost. For example, digital investment advisors can efficiently allocate an investor's portfolio across low-cost, passive ETFs,¹⁶ and financial advisory services can help consumers identify how to optimize paying down debts in order to minimize interest expenses. All of these services rely on analyzing broad sets of financial data in order to improve investor and consumer choice and outcomes. The development of GenAI and Agentic AI tools are expanding access to such services, and hold promise in providing even more insightful and actionable recommendations, as well as automated execution.

Another promising area of AI development in financial services involves a broad range of compliance, reporting, security and fraud detection, and trade and transaction monitoring functions (collectively referred to as "RegTech"). As noted above, a key strength of AI is its ability to ingest large volumes of data and identify patterns, anomalies, and insights that traditional approaches would be incapable of processing or detecting. AI can accordingly improve the efficiency, effectiveness, and capabilities of financial institutions and their regulators to the benefit of consumers, law enforcement, and market integrity. GenAI and Agentic AI are rapidly expanding RegTech capabilities, which as discussed further below in more detail, will require financial regulators to keep pace and adopt such tools themselves.

Examples of RegTech applications in financial services include improved AML monitoring, enhancing trade surveillance capable of identifying fraudulent and manipulative practices, and various forms of risk analysis, including with respect to market behavior and even economic indicators. As discussed further below, regulators are also leveraging these technologies—which should be supported by Congress—to better supervise markets and regulated entities, as well as to keep pace with combatting new technology-enabled threats.¹⁷

¹⁵ Of course, as discussed further below, not all models are well-designed and the quality and fairness of an output will be heavily dependent on the quality of the data input.

¹⁶ Beketov, M., Lehmann, K., and Wittke, M. (2018) 'Robo Advisors: quantitative methods inside the robots', *Journal of Asset Management*, 19(6), pp. 363-370. Available at: <https://link.springer.com/article/10.1057/s41260-018-0092-9>

¹⁷ FINRA (2022) *Deep Learning: The Future of the Market Manipulation Surveillance Program*. Available at: <https://www.finra.org/media-center/finra-unscripted/deep-learning-market-surveillance>.

As with any area of technological innovation, there are important risks that AI technology poses. For example, AI risks include the potential for embedding and perpetuating bias, processing and training on poor quality data, failing to operate as expected, helping bad actors engage in fraudulent and illegal conduct, and driving herd behaviors.

More specifically, certain “black box” and other poorly designed models run the risk of inadvertently embedding or even reinforcing bias by failing to test or understand outcomes, search for less discriminatory alternatives, and/or properly consider the input data driving the model results. As we have seen, models can also misfire or fail to perform as expected, resulting in a potential market “flash crash,” and divergence from our broader values. And, models that are not subject to proper governance, ongoing testing, and controls can drift from their prior performance and begin yielding faulty or flawed results based on changes in the data or underlying model conditions. I will discuss how these risks can and are being mitigated in the sections below.

Additionally, bad actors are increasingly able to use technology to engage in more sophisticated forms of fraud and illegal conduct, including through ID and voice cloning.¹⁸ These developments warrant focused law enforcement attention, continued development of new tools to combat these new threat vectors, and open collaboration between government, regulators, and industry to share information and best practices. As has consistently been the case, the best way to fight constantly evolving forms of crime is to ensure that law enforcement, regulators, and firms have the resources, information, and tools to confront bad actors.

Despite risks, the mere speculative potential or fear of future harm should not broadly block or disincentivize development and adoption of AI and emerging technologies in financial services, including by those small firms and community banks seeking to remain competitive in an increasingly digital economy. Global competitors are actively investing in AI technologies given their potential to transform most sectors and operations of our economies. This is an area that the U.S. must invest and lead, not just for the overall economic competitiveness of the country, but for all Americans who demand and deserve access to the most sophisticated financial markets and services in the world.

II. Financial Services Laws and Regulations: Capable of Integrating AI Tools and Technologies.

As noted above, financial services law, regulations, and risk management frameworks have long governed the adoption of emerging technologies in the sector, whether through rules to ensure consumer protection or principles-based guidance to ensure safety and soundness. These laws,

¹⁸ Brown, S. et al (2023) *Letter from Chairman Brown and Senators Menendez, Smith and Reed to Director Rohit Chopra*, United States Senate Committee on Banking, Housing, and Urban Affairs. Available at: https://www.banking.senate.gov/imo/media/doc/voice_cloning_financial_scams.pdf

regulations, and risk management principles largely apply to conduct and activities regardless of the technological tools used by the regulated entity—in this way they are appropriately technology-neutral. It is accordingly appropriate to focus attention on how to apply these existing laws and regulations in a manner that clarifies regulatory expectations, protects consumers and markets, and *encourages* firms to continue developing and adopting AI tools that improve their performance, safety, efficiency, and customer outcomes.

More specifically, when it comes to consumer protection, important safeguards apply when firms use AI tools. In the consumer finance space, lenders are subject to the Equal Credit Opportunity Act (ECOA), the Truth in Lending Act (TILA), the Fair Housing Act (FHA), the Fair Credit Reporting Act (FCRA), and prohibitions against unfair and deceptive acts and practices, amongst others. ECOA and the FHA specifically bar discrimination in lending for consumer credit, including in the context of residential transactions, and FCRA ensures key protections regarding credit bureau data use and access.

In the context of AI-based underwriting models, lenders must comply with fair lending regulations and search for less discriminatory alternatives (LDA) if a model results in a disparate impact on certain populations. To this end, AI can help raise the bar when it comes to understanding how certain variables may correlate with protected class characteristics, how that reliance can be reduced, and how less discriminatory alternative models can be deployed. Of course, not all models are well-built and governed. Regulators have authority, however, to examine fair lending compliance, including when lenders adopt new models or partner with third-party non-bank technology providers.

In other financial services contexts, consumer protection regulations require financial firms to provide consumers with viable avenues for lodging complaints, seeking recourse, and receiving broader customer support. Regulators have previously emphasized that consumer protections already apply to AI technologies in the context of chatbots, noting that “[d]epending on the facts and circumstances, entities may be subject to liability under federal consumer financial laws when chatbots fail to meet relevant requirements.”¹⁹ An example may be if a chatbot precludes a customer from being able to report and seek recourse on a fraudulent or incorrect payment transaction.²⁰

Beyond the consumer finance space, important safeguards also apply to applications of AI technologies in the capital markets. For example, following the 2010 flash crash, stronger protections were put in place by Congress to enforce against fraud and manipulation in trading

¹⁹ CFPB (2023) *Chatbots in consumer finance*. Available at: <https://www.consumerfinance.gov/data-research/research-reports/chatbots-in-consumer-finance/chatbots-in-consumer-finance/>.

²⁰ See, e.g., 12 CFR Part 1005 (Regulation E).

activity, including with respect to algorithms that may engage in wash trading or spoofing.²¹ Regulators and market participants have also adopted increasingly powerful trade surveillance technologies, often underpinned by AI technologies, to detect potential wrongdoing.²² Additionally, financial market regulators have updated and imposed new circuit breaker programs, systems compliance and integrity, and electronic trading risk rules to mitigate flash crash risks, trading volatility, and market disruptions.²³

In the investment advisory space, the SEC has existing conflicts of interest requirements in place that ensure that investment advice and recommendations made by broker-dealers or investment advisors are in the “best interest” of the investor.²⁴ Additionally, the SEC and FINRA have marketing rules that govern forms of investor advertising and engagement. These requirements apply equally to digital investment advisors, including those leveraging AI technologies, and as discussed further below have proven fully capable of providing the SEC with the tools it needs to ensure compliance and enforce against violations.²⁵

The above consumer and investor protection laws and regulations are by no means exhaustive, but are intended to give a broad overview of the existing frameworks that apply in financial services to adoption of AI and other emerging technologies. Additionally, beyond these consumer and investor protection laws, financial regulators have developed robust risk management principles and guidance to ensure the prudent, safe, and sound adoption of new technology-based models and systems.

More specifically, in addition to third-party and other risk management frameworks, financial institutions have long adhered to “Model Risk Management” (MRM) practices that apply to adoption of AI-based models. In 2011, the federal banking regulators released “SR 11-7: Guidance on Model Risk Management” (the “MRM Guidance”), which governs adoption of models and ways to mitigate associated risks. The MRM Guidance covers model design, documentation, governance, data use, performance, conceptual soundness, and ongoing monitoring, testing, and reporting considerations and expectations. The MRM Guidance is a

²¹ King & Spalding LLP (2019) *Spoofing: US Law and Enforcement*. Available at https://www.kslaw.com/attachments/000/007/109/original/Spoofing_US_Law_and_Enforcement.pdf?1564767398 (“In 2010, the Dodd-Frank Act amended the CEA to include spoofing as a disruptive practice. The anti-spoofing provision, CEA Section 4c(a)(5)(C), makes it unlawful for any person to engage in spoofing . . .”)

²² FINRA (2022) *Deep Learning: The Future of the Market Manipulation Surveillance Program*. Available at: <https://www.finra.org/media-center/finra-unscripted/deep-learning-market-surveillance>.

²³ SEC (2011) *Investor Alerts and Bulletins Investor Bulletin: New Stock-by-Stock Circuit Breakers*. Available at: <https://www.sec.gov/oiea/investor-alerts-bulletins/investor-alerts-circuitbreakers>; See also CFTC (2021) *Electronic Trading Risk Principles Final Rule*, 86 Fed. Reg. 2048; See also SEC (2020) *Staff Report on Algorithmic Trading in U.S. Capital Markets*. Available at: https://www.sec.gov/files/algo_trading_report_2020.pdf.

²⁴ SEC (2022) *Staff Bulletin: Standards of Conduct for Broker-Dealers and Investment Advisers Conflicts of Interest*. Available at: <https://www.sec.gov/im/iabd-staff-bulletin-conflicts-interest>.

²⁵ See, e.g., SEC (2021) *Schwab Subsidiaries Misled Robo-Adviser Clients about Absence of Hidden Fees*. Available at: <https://www.sec.gov/news/press-release/2022-104>.

primary framework governing adoption of models by financial institutions, including those developed by third parties.

In 2021, the OCC adopted updated MRM Guidance and specifically discussed its coverage of AI models; the SEC, FINRA, and the CFTC require similar risk management practices to be implemented by regulated entities.²⁶ Beyond imposition of general MRM frameworks, financial regulators will frequently issue regulations that govern related cybersecurity and system safeguards and integrity requirements, as well as guidance on issue-specific areas, including in the context of AI development and adoption.

Importantly, in the context of criminal conduct and scams that may also impact consumers, criminal laws are well-established that allow law enforcement and regulators to pursue bad actors. Additionally, law enforcement and intelligence agencies are actively investigating the use of new technologies to engage in criminal conduct and offering constructive guidance to the industry. In recent years, the NSA, FBI, and CISA published a report titled “Contextualizing Deepfake Threats to Organizations,” in which these agencies recommended to organizations that they pursue a multi-prong strategy of adopting new technologies to defend against bad actors, engage in information sharing, train personnel, and conduct scenario and contingency planning exercises.²⁷

III. Policy Recommendations to Ensure U.S. Leadership in AI Technologies that Can Advance and Improve Financial Markets and Services.

The above discussion makes clear that advances in the use of AI in financial services are developing within robust financial services legal and regulatory frameworks generally designed to mitigate risks associated with the use of emerging technologies, while enabling innovation that will benefit consumers, investors, economic dynamism, compliance, and U.S. competitiveness. That said, this discussion also makes clear that as with any area of technological advancement, we will need to evolve how we apply governing frameworks and be vigilant in identifying novel risks that might require tailored and specific regulatory interventions. The following are principles and recommendations that can help to strike the appropriate balance:

²⁶ FINRA (2020) *AI in the Securities Industry: Key Challenges and Regulatory Considerations*. Available at: <https://www.finra.org/rules-guidance/key-topics/fintech/report/artificial-intelligence-in-the-securities-industry/key-challenges>

²⁷ NSA, FBI, and CISA (2023) *Contextualizing Deepfake Threats to Organizations*. Available at: <https://media.defense.gov/2023/Sep/12/2003298925/-1/-1/0/CSI-DEEPFAKE-THREATS.PDF>

1. *Regulators Should Ensure that AI Adoption is not Blocked at the Examination Level and Should Use ‘Soft’ Power to Support Further Development and Applications.*

American financial institutions frequently report that innovation, including the adoption of AI technologies, is blocked by regulators at the examination level due to vague and endless inquiries and expectations that lack clear paths to compliance or conclusion.²⁸ Federal Reserve Vice Chair for Supervision Bowman recently highlighted the risk of regulators using the “soft” power of supervision to discourage [the] use” of new technologies.²⁹ As discussed further below, examiners should accordingly be well-versed in the benefits and risks of AI technologies, but specifically directed to avoid behaviors that unnecessarily chill adoption. The “tone at the top” of financial regulators must similarly create clear expectations for examination teams and market participants, encouraging appropriate safeguards without deterring further AI development and use. As Vice Chair Bowman aptly noted, a regulator’s soft power plays a significant role in creating an environment that fosters and advances responsible innovation.

With respect to rulemakings, since existing regulatory frameworks are largely fit-for-purpose in identifying and mitigating risks associated with AI technologies, it is important for regulators to avoid hasty and prescriptive rules based on speculative or hypothetical future risks. Instead, as discussed in more detail in the next section, policymakers and regulators should work to understand AI-related advances, how they are actually being adopted in the sector, benefits they are generating, and potential new risks they may pose. Only once actual risks that threaten to materially undermine consumer protection or broader safety and soundness are identified should targeted policy or regulatory interventions be considered. An approach of scatter-shot preemptive policies will otherwise deter or box-in further innovation and adoption, increase costs for firms, especially small financial institutions and community banks, and squander the first-mover and competitive advantages U.S. companies currently hold.

An example of a regulatory proposal that was overbroad and premature in its requirements related to AI-technologies was the SEC’s proposed rulemaking regarding “Conflicts of Interest Associated with the Use of Predictive Data Analytics by Broker Dealers and Investment Advisers” during the past Biden Administration. The proposed rule offered an overbroad definition of covered technologies, including AI technologies that have long been in existence, implied and asserted that such technologies pose a greater risk of harm to investors and going undetected (despite clear historical evidence

²⁸ Bank Policy Institute (BPI) (2025) *Request for Information on the Development of an Artificial Intelligence (AI) Action Plan*. Available at: <https://bpi.com/wp-content/uploads/2025/03/BPI-AI-Action-Plan-RFI-Mar-2025.pdf>.

²⁹ Bowman, M. (2025) ‘Bank Regulation in 2025 and Beyond’, *Speech by Governor Bowman at the Kansas Bankers Association Government Relations Conference, Federal Reserve*. Available at: <https://www.federalreserve.gov/newsevents/speech/bowman20250205a.htm>.

of misconduct made harder to trace by legacy tools, such as unscrupulous brokers targeting investors by telephone),³⁰ and then sought to treat firms using these technologies differently from those relying on legacy systems by subjecting them to a more punitive conflict of interest framework.

By specifically targeting emerging technologies largely based on conduct already subject to regulation or speculative conduct that may (or may not) occur in the future, the proposal was inherently not technology-neutral and would have deterred adoption of such technologies, especially for smaller firms that lack large compliance teams capable of parsing ambiguous regulatory expectations.³¹ Fortunately, this proposed rule has been withdrawn.

2. *Financial Regulators Require the Expertise, Tools, and Culture to Keep Pace with Increasingly Digital and Technology-Driven Financial Markets and Services.*

I had the distinct honor of serving as the first U.S. CFTC Chief Innovation Officer and Director of LabCFTC, the Agency’s bipartisan financial technology innovation initiative that became an office within the Agency. The mission of LabCFTC was to “facilitate market-enhancing innovation, inform policy, and ensure that the Agency [had] the technological and regulatory tools and understanding to keep pace with changes to our markets.” Efforts like LabCFTC should not be mere public relations efforts, but rather should play a critical and substantive external and internal role within an Agency.

More specifically, innovation offices can serve as an external door to the public—and an entry point where the public, startups, incumbents, and other interested stakeholders alike can engage with an Agency. Through engagement—and proper coordination with Agency operating divisions and staff—an innovation office can help an Agency keep up with a rapidly changing external environment. Such an office should work across the Agency to be responsive to the marketplace, offering clear expectations when ambiguity is identified or helping to inform rulemakings. It should further help to internalize learnings to improve the Agency’s overall understanding of new technologies, including AI.

This is not a regulatory race to the bottom—it is the opposite: it is a means to modernize and support regulatory excellence. A recent GAO report underscores this point given its finding of a significant lack of relevant technology skills and understanding across

³⁰ SEC (no date) *Boiler room schemes*. Available at: <https://www.investor.gov/introduction-investing/investing-basics/glossary/boiler-room-schemes>.

³¹ Gorfine (2019)

federal financial regulators.³² A better informed Agency can perform all functions more effectively, ranging from avoiding scenarios where examiners block regulated entities because they simply don't know or understand whether a new technology or application is appropriate to scenarios where enforcement teams miss new categories of crime due to a lack of awareness.

In addition, innovation offices can be key hubs focused on attracting the best and brightest in a technological field and can help an Agency acquire leading-edge technologies. Indeed, in a world where markets and services are increasingly digital, quantitative, and technology-driven, a regulator must be increasingly digital, quantitative, and technology-driven as well. Only by adopting leading technologies, including AI, can regulators properly regulate, surveil, and monitor their proper jurisdictions.³³ AI tools will prove increasingly critical in combatting fraud, ensuring market integrity, analyzing financial information—including real-time information about the economy or even individual firms—and identifying concentrations of risk.

I applaud this Subcommittee for past bipartisan legislative efforts to codify innovation offices like LabCFTC, and encourage getting these efforts across the finish line. American taxpayers deserve financial regulators built for the second half of the 21st century, and Agency modernization, which must include retaining critical internal subject matter expertise and technology capabilities, should be a top priority.

3. Establish a Federal Data Privacy Framework that Contemplates the Evolution of AI and Ensure Open Access to Quality, Permissioned Financial Data.

As previously discussed, the quality of AI model outputs and predictions is inherently tied to the quality of data inputs used to train and operate such models. Policymakers, regulators, and the broader public are right to be concerned about whether we have appropriate safeguards in place regarding how data is being generated, sourced, secured, shared and consumed for use in AI models.³⁴ To this end, it is imperative that Congress work to establish a national framework that governs data privacy, advances cybersecurity, and ensures that consumers have control over—and trust and confidence in—how their data and information is being used.

³² GAO (2023) 'FINANCIAL TECHNOLOGY: Agencies Can Better Support Workforce Expertise and Measure the Performance of Innovation Offices', *Report to the Chairman, Committee on Financial Services, House of Representatives*. Available at: <https://www.gao.gov/assets/gao-23-106168.pdf>

³³ Giancarlo, J.C. (2018) 'Quantitative Regulation: Effective Market Regulation in a Digital Era,' *Keynote Address of Chairman J. Christopher Giancarlo at Fintech Week, Georgetown University Law School*. Available at: <https://www.cftc.gov/PressRoom/SpeechesTestimony/opagiancarlo59>.

³⁴ Zakrzewski, C., Lima, C., and DiMolfetta, D. (2023) *Tech leaders including Musk, Zuckerberg call for government action on AI*. Available at: <https://www.washingtonpost.com/technology/2023/09/13/senate-ai-hearing-musk-zuckerburg-schumer/>.

The Gramm-Leach-Bliley Act (GLBA) provides a baseline for how a broad range of financial firms must explain their information-sharing practices and safeguard consumer data. This law may provide a good starting point for the regulation of data use and security in financial services, but is due for modernization given market developments, including with respect to AI and the increasing role played by nonbanks and technology providers. Key areas for legislative updating include increased consumer transparency and control over the use of personal data, implementation of data minimization principles so that data is used for clear and stated purposes, opt-out rights for nonessential data collection, broader coverage, and consumer rights to withdraw consents or delete personal data, as appropriate.

On the latter point of data deletion rights, it will be important for policymakers to consider how a consumer's invocation of this right could impact his or her future access to certain services, including lending or insurance, that may be limited by a lack of data access or availability. The law will need to balance the notion that a consumer should not ever be penalized or targeted for exercising a data control right with the reality that future underwriting models may require (or prefer) access to certain types of data. For this reason, consumers need to be fully informed of the implications in electing certain data treatment, including potential downstream effects in securing future services.

As discussed further below, it is also important to note that a modernized national framework over data privacy and use should preempt overlapping state-based efforts. Our economy is increasingly digital, which inherently means that Internet and mobile-based platforms will increasingly be leveraged in the financial services context. These platforms (and the data moving on them) operate on a national (if not global) basis, which renders state-based frameworks inefficient, costly, confusing, and potentially unworkable. A patchwork approach to data regulation may accordingly undermine AI development in the U.S., result in market fragmentation, and also result in consumers, small and community banks, and other market participants in certain states being excluded from the benefits of such innovation.

As a final matter relating to the centrality of data in AI development, it is important that the CFPB confirm a robust open banking rule that allows consumers the ability to share their data with other financial services providers. Open banking is predicated on the notion that consumers have a right to control the sharing and use of their financial data, and that such data should not be subject to anti-competitive restrictions on its transfer. In addition to upholding these principles, the CFPB should amend the prior Administration's final rule that imposed overbroad limitations on secondary use of permissioned consumer

data since it unnecessarily treats open banking data differently from all other data and will impede AI model development, including in the context of countering fraud.

Overall, GenAI and Agentic AI are still in their early stages of development and in many applications will heavily rely on the ingestion of high-quality, real-time, and permissioned consumer data to unlock consumer benefits. For example, Agentic AI will increasingly be able to offer a consumer constant financial health monitoring based on real-time information and be able to recommend and execute optimal financial decisions, ranging from paying off high-cost debt to shifting savings to a high-yield account.³⁵ This requires the seamless flow of high-quality data subject to appropriate safeguards. For this reason, national data privacy, open banking, and data security frameworks in the U.S. must be crafted with the future of AI development in mind.

4. *Prevent the Interference of State Laws and Regulations with the Federal Framework for Regulating Financial Services.*

As detailed above, there is a longstanding and robust federal framework for the regulation of financial services and markets in the U.S.—a framework that contemplates a role for state banking regulators and that has successfully allowed for the integration of emerging technologies over the decades. Indeed, the financial industry should be looked at as an example for how other sectors can safely adopt emerging technologies, including AI.

Unfortunately, at the state level, there have been numerous legislative efforts in recent years to regulate AI technologies or the data they require without careful consideration of the interplay or impact on overarching financial regulation.³⁶ This raises significant risks of conflict, ambiguity, redundancy, and confusion—any one of which will impede AI development and adoption in financial services. It is accordingly important that Congress act at the federal level to ensure a consistent and unified national approach given the inherently national nature of AI development and use in financial services. Failure to act will inevitably chill adoption and undermine compliance efforts across the industry.

³⁵ Zhang, B. and Garvey, K. (2025) 'From automation to autonomy: the agentic AI era of financial services', University of Cambridge, Judge Business School. Available at: <https://www.jbs.cam.ac.uk/2025/from-automation-to-autonomy-the-agentic-ai-era-of-financial-services/>

³⁶ Welle, L.J. et al. (2025) *The Evolving Landscape of AI Regulation in Financial Services*. Available at: <https://www.goodwinlaw.com/en/insights/publications/2025/06/alerts-finance-fs-the-evolving-landscape-of-ai-regulation>

5. *Further Clarify and Update Risk Management Frameworks and Promote the Value of Well-Crafted Standards through Regulatory Recognition and Safe Harbors.*

As discussed above, the financial services industry is subject to longstanding risk mitigation frameworks that have been applied to the adoption of emerging technologies for decades. These frameworks are appropriate for advances in AI in financial services, but will require updates, clarifications, and further guidance that address emerging risks and provide the industry with the clarity required to confidently build and adopt such technologies.

For example, it is useful and appropriate for regulators to alert and reinforce that the use of AI in customer service applications must provide customers with proper recourse and paths to resolution; such notice should underscore that there is a tremendous opportunity to improve on legacy customer service models, which historically have included long waits for accessing live representatives, painful dropped calls or endless routing loops, and failure to achieve proper resolution. As financial providers develop, test, and gradually implement new AI-powered models with greater functionality that can improve on the status quo, it would be appropriate for regulators to emphasize the importance of risk mitigants, such as ensuring proper audits and performance monitoring of such models.

Other key areas for clarifying or enhancing guidance and regulatory expectations in the context of AI models include MRM Guidance and third-party risk management with a focus on specific consumer finance and financial market applications. For example, while the banking MRM Guidance and 2021 OCC updates noted above provide the appropriate framework for assessing and mitigating risk with AI models, regulators should look for opportunities to clarify how the Guidance should be applied in the context of AI. Confirmation that establishes that the mere use of AI—even GenAI—does not inherently make an activity higher-risk would be especially valuable, including for smaller financial institutions.

Additionally, regulators should look for opportunities to make clear to the market which specific applications should be viewed as low, moderate, or high-risk in order to help providers tailor compliance requirements. To this end, regulators should look for opportunities to make clear that degrees of explainability that are required will vary depending on the risk of the application and availability of other risk mitigants, including ongoing performance testing or the presence of a “human in the loop.”³⁷ And in the financial markets context, market regulators should further clarify expectations for

³⁷ Google Cloud and AIR (2024) *Adapting Model Risk Management for Financial Institutions in the Generative AI Era*. Available at: <https://regulationinnovation.org/adapting-model-risk-management-for-financial-institutions-in-the-generative-ai-era/>

exchanges and trading firms relying on AI technologies, including proper risk controls that can mitigate herding behavior, new forms of manipulation, or excessive trading volatility caused by technology failures.

Finally, policymakers and regulators should encourage and help foster the development of standards that can enhance safety, security, and adoption of AI technologies. The National Institute of Standards and Technology (NIST) previously released its “Artificial Intelligence Risk Management Framework 1.0,” which provides for approaches to mitigating risks associated with the use of AI.³⁸ Private sector and industry organizations often build on NIST and similar standards efforts to create even more tailored frameworks that can help ensure compliance in the financial services context. For example, the [Cyber Risk Institute \(CRI\)](#) Profile was developed by a “not-for-profit coalition of financial institutions and trade associations working to protect the global economy by enhancing cybersecurity through standardization. Through consensus among the financial sector ecosystem... [the Profile] and related guidance help firms better manage cyber compliance programs.”

Such industry-led efforts to develop standards in the AI context should be strongly supported by regulators. Indeed, I recently argued that regulators will inevitably need to rely on external standards development given the speed of technological change and the lack of internal capabilities to keep pace with all technology-driven developments.³⁹ To this end, regulators can advance standards development—including in the context of AI—by formally recognizing standards setting organizations and crafting explicit compliance safe harbors that encourage adherence to such standards.

*

*

*

As detailed in my testimony, the development and adoption of AI in financial services is not only inevitable, but it holds tremendous promise. This technology can improve the delivery, accessibility, and cost of financial services, while also better serving consumers and businesses. This evolution is also vital to ensuring the global competitiveness of our financial firms and the enduring appeal of U.S. financial markets in a world of increasing competition.

While it is wholly appropriate for regulators to closely monitor these developments and target interventions when specific risks are identified, it is equally important to enhance clarity and encourage the industry to continue developing and responsibly adopting AI. Given existing financial services regulatory frameworks, the adoption of complex models in higher-risk

³⁸ Crossman, P. (2023) *How banks should deal with AI's risks, according to NIST*. Available at: <https://www.americanbanker.com/news/how-banks-should-deal-with-ais-risks-according-to-nist>.

³⁹ Gorfine, D. and Bailey, N. (2025) ‘The Case for Recognition of Technology Standards in Financial Regulation,’ *Open Banker*. Available at: <https://openbanker.bechiiv.com/p/techstandards>.

functions will naturally be gradual and subject to robust governance requirements. We should avoid creating redundancies, conflicts, or ambiguity that could inadvertently chill innovation, especially among small firms and community banks.

AI will be at the center of the second half of the 21st century. It is a technology that the U.S. must lead in, and by carefully balancing a focus on safeguards with a clear pro-innovation philosophy, we can ensure that American ingenuity thrives.

Thank you. I am happy to answer any questions that you have.

Appendix A*Daniel Gorfine Biography*

My name is Daniel Gorfine, and I am the founder and CEO of Gattaca Horizons LLC, a boutique advisory firm. I am also a co-founder and director of the non-profit Digital Dollar Project and an adjunct professor at the Georgetown University Law Center. I am honored to have previously served as Chief Innovation Officer at the U.S. Commodity Futures Trading Commission (CFTC) and Director of LabCFTC.

In my advisory capacity, I work with a range of clients, including financial firms, technology companies, fintech and digital asset-related firms, industry trade associations, and startups. As an adjunct professor at Georgetown, I teach a course titled “Fintech law and policy.” I am a graduate of Brown University (A.B.), hold a J.D. from the George Washington University Law School and an M.A. from the Paul H. Nitze School for Advanced International Studies (SAIS) at Johns Hopkins University.

Chairman STEIL. Thank you very much.
We now recognize Nicol Turner Lee for 5 minutes.

STATEMENT OF NICOL TURNER LEE, SENIOR FELLOW AND DIRECTOR, CENTER FOR TECHNOLOGY INNOVATION, BROOKINGS INSTITUTION

Ms. TURNER LEE. Thank you, Chairman Steil, Ranking Member Lynch, and distinguished members of the subcommittee. Thank you for this invitation to testify on the use of AI in the financial sector.

My name is Dr. Nicol Turner Lee. I am a director of the Center for Technology Innovation at the Brookings Institution, and I am also the co-author of a report that was done on AI in the global markets by the CFTC.

Artificial intelligence brings a variety of opportunities to the financial sector, and for years it has been used in banking, fraud detection, mortgage applications, credit underwriting, and data analytics. Many of these use cases are promising, offering the potential for greater accessibility and improved customer service, while others may introduce challenges, including concerns about racial bias and discrimination.

Given that the financial services industry is one of the country's most highly regulated ones, adoption and use of AI requires careful review. Without such oversight, risks abound that can undermine consumers' ability to be economically resilient, especially under changing economic conditions. Further, inaccurate or discriminatory information that is used to train AI models can put marginalized populations at even greater risk, threatening to widen the racial wealth gap, limit access to home ownership, credit, and other financial transactions. In other words, when algorithms make poor decisions, the quality of life for Black and Brown communities, seniors, and even some of us are placed in reverse.

For these and other reasons, Congress must continue to foster responsible and ethical use in the financial regulation sector by providing safeguards that protect consumers from AI risks now and reinforce algorithmic accountability, safety, and security, especially when we look at incidents of fraud.

Congress can also safeguard consumers from both the unintended and intended consequences through legislation, which starts with comprehensive national privacy standards and reinforcing the jurisdiction of independent Federal agencies and State attorneys general, who enforce consumer protection regulations in the age of AI.

For example, the former director of the Consumer Financial Protection Bureau clarified that algorithmic decisionmaking is held to the same standards as human decisionmaking. The Department of Justice, the Department of Housing and Urban Development, these agencies also previously released documents stating the compliance with many of the provisions that apply to either tenant screening or fairness overall.

But recent actions taken to defund and compromise the independence of Federal agencies like the CFPB, like the Federal Trade Commission (FTC), which have oversight over deceptive consumer practices, can have major implications on protecting consumers from the opaque technology of AI.

State and State attorneys general are critical in enforcing those protections, but challenges to their State rights and laws are not productive. Even with the rejection of a 10-year moratorium on States' ability to develop their own laws, the current AI Action Plan still suggests that AI-related funds should not go to those with burdensome AI regulations.

In the absence of a national framework, States must continue to protect the millions of seniors, children, marginalized populations, and even farmers who are being barraged by algorithmic discrimination, deepfakes, and other malicious attacks on their financial well-being.

Recent proposals from legislators seek to establish regulatory sandboxes for AI developers where companies can apply for waivers or modifications but without strong regulatory or enforcement consequences, companies will have another avenue to exploit the personal information and behaviors of consumers.

As my Brookings colleague Aaron Klein suggests, we should be creating more greenhouses, which allow for more transparency and sunlight into these processes, promoting collaborative partnerships between business, government, and consumers, to see where issues arise and how the three of them can foster trust with one another. These approaches can also assist in anti-bias discrimination, which is also very much imperative to building trust.

I would like to just close with five proposals for members of this committee to consider as we embark on the inquiry today.

Ensure that the industry is compliant with existing legal statutes and remedies which reinforce algorithmic accountability. The financial sector is already regulated by numerous Federal and State organizations. It is imperative that it stay that way, and that they also comply with acts like the Fair Credit Reporting Act and others to seek the prevention of discriminatory results.

Mandate transparency guidelines and full disclosure for all consumers. Companies need to tell people publicly when they are using AI to make decisions. Those matter outputs should be explainable and accurate and reliable.

Encourage, not discourage, responsible and ethical development of financial models. Innovation and regulation can be complementary, and I think that is to the benefit, again, of those greenhouses, which show transparency in the ideas and processes that we embrace in this area.

Brace for the adoption of agentic AI but do the first two. Make sure there is consumer protection in place that we are poised to benefit from the autonomous nature of agentic AI, but we still have firm oversight.

Finally I would say, invest in AI financial literacy programs so that consumers know how this industry and sector is evolving.

These and other questions I place before this committee, and I look forward to your questions and collaboration.

[The prepared statement of Ms. Turner Lee follows:]

Written Testimony of Nicol Turner Lee
Director, Center for Technology Innovation
Senior Fellow, Governance Studies
The Brookings Institution

Before the United States House of Representatives Committee on Financial Services
Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence
Hearing on “Unlocking the Next Generation of AI in the U.S. Financial System for
Consumers, Businesses, and Competitiveness”

September 18, 2025

Chairman Steil, Ranking Member Lynch, and distinguished members of the Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence: thank you for the invitation to testify on the use of artificial intelligence in the financial sector. I am Nicol Turner Lee, a Senior Fellow in the Governance Studies program and Director of the Center for Technology Innovation at the Brookings Institution. With a history of over 100 years, Brookings is committed to evidence-based, nonpartisan research in a range of focus areas. My research expertise encompasses data collection and analysis around regulatory and legislative policies that govern telecommunications and high-tech industries, along with the impacts of innovation, artificial intelligence, and the digital divide. I am also the author of *Digitally Invisible: How the Internet is Creating the New Underclass*, which explores and investigates the latter point. I am testifying in my individual capacity.

I. Introduction

The financial sector has long adopted many of the newest technologies to increase efficiency and profits while minimizing risk. Artificial intelligence (AI) brings opportunities to advance these goals, especially to support information processing, and was quickly adopted by the retail financial

services sector. In fact, the sector's spending on AI is expected to grow to \$97 billion by 2027, up from \$35 billion in 2023.¹ For years, AI has been used in banking, fraud detection, mortgage applications, and credit underwriting in addition to other backend financial functions like data analytics. Many of these use cases are promising, offering the potential for greater accessibility and improved customer service, while others may introduce challenges, including concerns about bias and discrimination.

Yet the technology was introduced into a complex regulatory landscape. As one of the country's most highly regulated industries, there are several independent agencies with oversight over financial institutions, including the Federal Deposit Insurance Corporation (FDIC), Securities and Exchange Commission (SEC), Consumer Financial Protection Bureau (CFPB), and Commodity Futures Trading Commission (CFTC).² Each of these is designated unique or overlapping authority to ensure regulatory compliance to achieve goals such as financial stability and taxpayer protection. The expansion of AI into the sector also has called for guidance on how these regulations apply to the technology and an assessment of where additional safeguards are needed. Such efforts progressed in part under the prior administration.

In 2023, the CFPB released guidance on credit denials involving AI, describing how "lenders must use specific and accurate reasons when taking adverse actions against consumers."³ Seeking to empower the directors of the Federal Housing Finance Agency (FHFA) and CFPB, additional guidance in 2023 required regulated entities "to use appropriate methodologies including AI tools to

¹ Parker, David. "The Future of AI in Financial Services." *Forbes*, October 3, 2024.

<https://www.forbes.com/sites/davidparker/2024/10/03/the-future-of-ai-in-financial-services/>.

² Labonte, Marc. "Introduction to Financial Services: The Regulatory Framework." Congressional Research Service, April 1, 2025. <https://www.congress.gov/crs-product/IF11065>.

³ Consumer Financial Protection Bureau. "CFPB Issues Guidance on Credit Denials by Lenders Using Artificial Intelligence," September 19, 2023. <https://www.consumerfinance.gov/about-us/newsroom/cfpb-issues-guidance-on-credit-denials-by-lenders-using-artificial-intelligence/>.

ensure compliance with Federal law,” “evaluate their underwriting models for bias,” “evaluate automated collateral-valuation and appraisal processes in ways that minimize bias,” and “combat unlawful discrimination enabled by automated or algorithmic tools used to make decisions about access to housing and in other real estate-related transactions.”⁴ While these measures are critical as the use of AI continues to grow in the industry, they have since been revoked with work stopped or eliminated at the current downsized CFPB.⁵

In this written testimony, I argue that the adoption and use of AI in the financial services sector are at an accelerated pace and require careful review. Risks abound that can undermine consumers’ ability to be economically resilient and experience the benefits of wealth creation. Due to inaccurate or discriminatory information that is often used to train AI models, marginalized populations are at even higher risk when AI systems make misjudgments about their qualifications and eligibility for economic opportunities; thereby, widening the wealth gap and limiting access to homeownership, credit worthiness, and other financial transactions that bring hope and promise to their quality of life. For these reasons, Congress must continue to foster responsible and ethical AI use in the financial regulation sector by providing safeguards that protect consumers from AI risks now and into the future.

II. AI in financial services requires responsible and ethical deployment and design to protect vulnerable populations

⁴ Biden, Joseph R. “Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence.” The White House, October 30, 2023. <https://bidenwhitehouse.archives.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>.

⁵ Rugaber, Christopher. “Trump Administration Orders Consumer Protection Agency to Stop Work, Closes Building.” *Associated Press*, February 9, 2025. <https://apnews.com/article/trump-consumer-protection-cease-1b93c60a773b6b5ee629e769ae6850e9>.

Across all sectors, it is important that AI must be responsibly and ethically designed. This process begins early on, ensuring fairness and accuracy in the data used to train models, consistently auditing outputs, and identifying high-risk use cases. Many uses of AI in the financial sector are considered high impact, as credit score calculations or loan denials can have major consequences on individuals' lives and financial opportunities.⁶ These risks are not abstract or hypothetical; research has continued to demonstrate discriminatory results using AI in decision-making.

While human decision-making in the financial sector undoubtedly introduces some bias and AI may show promise in eliminating such inequities through potential technical objectivity, at its current stage, research shows some models reinforce them. For example, the data used to calculate eligibility and credit scores often reflect historical inequities. Black and Hispanic communities are more likely to have limited credit history and may not be adequately represented in the data used to train a model to carry out financial eligibility decisions.⁷ Put into context: A 2023 report from Pew Research Center shows that white and Asian households reported significantly more wealth than Black, Hispanic, and multiracial counterparts.⁸ The same research showed that just 45% of Black households were categorized in the middle or upper wealth tiers in 2021, compared to over 70% of white and Asian ones.⁹ Whereas past discrimination such as redlining may have directly correlated with one's race, lending applications can use location, names, or other information to signal an

⁶ Klein, Aaron. "Credit Denial in the Age of AI." Brookings, April 11, 2021. <https://www.brookings.edu/articles/credit-denial-in-the-age-of-ai/>; Valdrighi, Giovanni, Athyrson M. Ribeiro, Jansen S. B. Pereira, Vitoria Guardieiro, Arthur Hendricks, Décio Miranda Filho, Juan David Nieto Garcia, et al. "Best Practices for Responsible Machine Learning in Credit Scoring." *Neural Computing and Applications* 37 (August 5, 2025): 20781–821. <https://doi.org/10.1007/s00521-025-11520-y>.

⁷ Munsterman, Korin. "When Algorithms Judge Your Credit: Understanding AI Bias in Lending Decisions." Accessible Law. UNT Dallas Law Review, 2025. <https://www.accessiblelaw.untdallas.edu/post/when-algorithms-judge-your-credit-understanding-ai-bias-in-lending-decisions>.

⁸ Kochhar, Rakesh, and Mohamad Moslimani. "Wealth Gaps across Racial and Ethnic Groups." Pew Research Center, December 4, 2023. <https://www.pewresearch.org/2023/12/04/wealth-gaps-across-racial-and-ethnic-groups/>.

⁹ Ibid.

applicant's race, even if it's not explicitly reported.¹⁰ When AI models are trained on discriminatory or inaccurate data, it further harms and often disqualifies consumers who have been trapped by structural inequalities that limit wealth creation.

Extending beyond AI's impact on historically marginalized communities, emerging technologies also enable financial institutions to rely on a wider range of data sources to make these decisions, many of which are not directly related to financial status or that may be used as proxies for such. Gross inequities can arise from proxy measures deployed by financial institutions. For example, information such as what type of device a user is on, their email provider, whether someone capitalizes the first letters of their name, and even social media histories can be used to predict financial behaviors.¹¹ These decisions extend beyond whether someone receives a loan; such scores may also be considered by landlords when considering a tenant or by companies when making a hiring decision.¹²

III. Financial safeguards must be in place to ensure trustworthy AI

Congress must focus on consumer protection to safeguard consumers from the aforementioned biases and the unsubstantiated decisions that may arise when automated decision-making systems are used. Everyday consumers should be provided with assurance via financial disclosures and enforcement against bad actors. Just as AI tools could be used by businesses to gain efficiency and analyze greater swaths of information, they are also accessible to scammers and cyber criminals for the same ends. Scammers can use generative models to replicate the voice or likeness

¹⁰ Bowen III, Donald E., S. McKay Price, Luke C.D. Stein, and Ke Yang. "Measuring and Mitigating Racial Bias in Large Language Model Mortgage Underwriting." *SSRN Electronic Journal*, April 30, 2024. <https://doi.org/10.2139/ssrn.4812158>.

¹¹ Klein, Aaron. "Reducing Bias in AI-Based Financial Services." Brookings, July 10, 2020. <https://www.brookings.edu/articles/reducing-bias-in-ai-based-financial-services/>.

¹² TechEquity. "Unpacking HUD's New Guidance on Algorithmic Tenant Screening." May 29, 2024. <https://techequity.us/2024/05/29/unpacking-huds-new-guidance-on-algorithmic-tenant-screening/>.

of someone a target might trust, or better gather information on individuals, to create more effective social engineering schemes.¹³ A report from the Identity Theft Research Center recently found that impersonation has become the most reported type of scam, with incidents increasing 148% from 2024-2025.¹⁴ Older consumers may be more likely to fall prey to these deepfakes and lose their life savings.¹⁵ Such thefts are devastating for senior citizens and can completely ruin their lives.

The role of independent agencies and states

Independent federal agencies and state attorneys general play a large role in holding AI models accountable, as well as big business. In the absence of comprehensive federal regulation over AI's use in financial services, the Federal Trade Commission (FTC), CFPB, Department of Housing and Urban Development (HUD), and state attorneys general have extended existing consumer protection regulations and frameworks to the AI context. For example, the CFPB previously published research on the use of automated customer service models in consumer finance.¹⁶ The former CFPB Director Rohit Chopra also stated that "there is no 'fancy new technology' carveout to existing laws," emphasizing that algorithmic decision-making is held to the same standards as human decision-making.¹⁷ Since courts have ruled that firms' use of biased algorithmic tools may constitute

¹³ Gedy, Grace. "AI Voice Cloning: Do These 6 Companies Do Enough to Prevent Misuse?" Consumer Reports, March 10, 2025. <https://innovation.consumerreports.org/AI-Voice-Cloning-Report-pdf>; Park, Peter S., Simon Goldstein, Aidan O'Gara, Michael Chen, and Dan Hendrycks. "AI Deception: A Survey of Examples, Risks, and Potential Solutions." *Patterns* 5, no. 5 (May 10, 2024). <https://doi.org/10.1016/j.patter.2024.100988>.

¹⁴ Identity Theft Resource Center. "2025 Trends in Identity Report," June 2025. <https://www.idtheftcenter.org/wp-content/uploads/2025/06/2025-ITRC-Trends-in-Identity-Report.pdf>.

¹⁵ Bradley, Steven. "Behind the Screen: Elder Financial and Technology Abuse in the Age of AI." American Bar Association, June 24, 2025. https://www.americanbar.org/groups/law_aging/publications/bifocal/vol46/vol46issue5/elderabuseandartificialintelligence/; iProov. "iProov Study Reveals Deepfake Blindspot: Only 0.1% of People Can Accurately Detect AI-Generated Deepfakes," February 12, 2025. <https://www.iproov.com/press/study-reveals-deepfake-blindspot-detect-ai-generated-content>.

¹⁶ Consumer Financial Protection Bureau. "Chatbots in Consumer Finance," June 6, 2023. <https://www.consumerfinance.gov/data-research/research-reports/chatbots-in-consumer-finance/chatbots-in-consumer-finance/>.

¹⁷ Consumer Financial Protection Bureau. "CFPB Comment on Request for Information on Uses, Opportunities, and Risks of Artificial Intelligence in the Financial Services Sector," August 12, 2024. <https://www.consumerfinance.gov/about-us/newsroom/cfpb-comment-on-request-for-information-on-uses->

a biased policy decision, the CFPB should continue to monitor firms' compliance with the Equal Credit Opportunity Act and the Consumer Financial Protection Act. In cases where AI is used for "fraud screening," firms may also have to comply with the Fair Credit Reporting Act.¹⁸ As early as 2023, the CFPB also issued guidance stating that entities using AI credit underwriting models to make lending decisions must be able to articulate "accurate" and "specific" reasons for adverse actions—they cannot hide behind the complexity of black-box models.¹⁹ The Department of Justice and HUD also previously released documents stating that Fair Housing Act (FHA) provisions apply to tenant screening algorithms.²⁰

However, the recent actions taken to defund and compromise the independence of federal agencies, such as the CFPB, can have major implications on protecting consumers from the use of opaque and discriminatory technology.²¹ That is why Congress must ensure that these independent agencies be directed to work for the good of the public before AI and Big Tech companies define what consumer protection is for everyday people who depend on their financial assets to thrive and survive in society.

State attorneys general are also applying existing consumer protections and state laws to AI.²² For example, the Massachusetts attorney general recently reached a million-dollar settlement with a

[opportunities-and-risks-of-artificial-intelligence-in-the-financial-services-sector/](#); Brennan, J. Scott Babwah, Kevin Frazier, and Anna Viñals Musquera. "Are Existing Consumer Protections Enough for AI?" *Lawfare*, September 3, 2025. <https://www.lawfaremedia.org/article/are-existing-consumer-protections-enough-for-ai>.

¹⁸ *Ibid*

¹⁹ Consumer Financial Protection Bureau. "CFPB Issues Guidance on Credit Denials by Lenders Using Artificial Intelligence," September 19, 2023. <https://www.consumerfinance.gov/about-us/newsroom/cfpb-issues-guidance-on-credit-denials-by-lenders-using-artificial-intelligence/>.

²⁰ Brennan, J. Scott Babwah, Kevin Frazier, and Anna Viñals Musquera. "Are Existing Consumer Protections Enough for AI?" *Lawfare*, September 3, 2025. <https://www.lawfaremedia.org/article/are-existing-consumer-protections-enough-for-ai>; TechEquity. "Unpacking HUD's New Guidance on Algorithmic Tenant Screening," May 29, 2024. <https://techequity.us/2024/05/29/unpacking-huds-new-guidance-on-algorithmic-tenant-screening/>.

²¹ Consumer Reports Advocacy. "Senate Passes Budget Bill with Devastating Cut to CFPB's Funding," July 1, 2025. https://advocacy.consumerreports.org/press_release/senate-passes-budget-bill-with-devastating-cut-to-cfpbs-funding/.

²² Taylor, Ashley, Clayton Friedman, and Gene Fishel. "State AGs Fill the AI Regulatory Void." *Reuters*, May 19, 2025. <https://www.reuters.com/legal/legalindustry/state-ags-fill-ai-regulatory-void-2025-05-19/>; Cook, Andrew. "State

company for failing to take reasonable measures to prevent fair lending risks in their AI underwriting models, which disproportionately penalized non-white and non-citizen applicants.²³ Some states, such as Oregon, New Jersey, and California, have also issued guidance confirming that AI systems may be subject to prosecution under consumer protection statutes preventing unfair or deceptive practices.²⁴ Outside of the financial services context, Texas Attorney General Ken Paxton recently opened an investigation into AI chatbot platforms for possible fraudulent claims, privacy misrepresentations, and other privacy and data security-related harms.²⁵ These moves by state officials will extend into financial services as more agentic AI is launched.

The caution of proposed regulatory sandboxes

But these state-led actions can be weakened by federal legislation that leans into more deregulation in the financial services sector. Recent attempts to establish sandboxes for AI developers where companies can apply for waivers or modifications to regulations are examples of such regress in progress.²⁶ Without regulatory or enforcement consequences, companies will have more reputational risks when such problems come to light during adoption stages. Rather, the integration of AI should be a collaborative partnership between businesses and consumers to see

Attorneys General on Applying Existing State Laws to AI.” Orrick, February 18, 2025.

<https://www.orrick.com/en/insights/2025/02/State-Attorneys-General-on-Applying-Existing-State-Laws-to-AI>.

²³ Massachusetts Office of the Attorney General. “AG Campbell Announces \$2.5 Million Settlement with Student Loan Lender for Unlawful Practices through AI Use, Other Consumer Protection Violations,” July 10, 2025.

<https://www.mass.gov/news/ag-campbell-announces-25-million-settlement-with-student-loan-lender-for-unlawful-practices-through-ai-use-other-consumer-protection-violations>

²⁴ Brennan, J. Scott Babwah, Kevin Frazier, and Anna Viñals Musquera. “Are Existing Consumer Protections Enough for AI?” *Lawfare*, September 3, 2025. <https://www.lawfaremedia.org/article/are-existing-consumer-protections-enough-for-ai>.

²⁵ Texas Office of the Attorney General. “Attorney General Ken Paxton Investigates Meta and Character.AI for Misleading Children with Deceptive AI-Generated Mental Health Services,” August 18, 2025.

<https://www.texasattorneygeneral.gov/news/releases/attorney-general-ken-paxton-investigates-meta-and-characterai-misleading-children-deceptive-ai>.

²⁶ U.S. Senate Committee on Commerce, Science, & Transportation. “Sen. Cruz Unveils AI Policy Framework to Strengthen American AI Leadership,” September 10, 2025. <https://www.commerce.senate.gov/2025/9/sen-cruz-unveils-ai-policy-framework-to-strengthen-american-ai-leadership>.

where issues arise and how the two can foster trust with one another. Regulatory sandboxes have previously been used for industries such as the financial technology (“fintech”) and pharmaceutical industries to generate evidence that can then advance policy and technological innovation.²⁷ However, they must be adopted with consumer protection safeguards in place, such as disclosure requirements and continuous monitoring. Regulatory sandboxes may also pave the way for anti-bias experimentation and early identification of potential risks and harms to consumers, within a controlled environment with safe harbor protections.²⁸ Rather than allowing a privileged set of firms to innovate at any cost while skirting accountability mechanisms, regulatory sandboxes should foster collaboration between regulators, developers, and consumers to co-create effective standards, and the process for arriving at consensus should be transparent to government, industry, and civil society actors.

IV. The AI Action Plan

Regulators and attorneys general have shown a keen interest in filling in gaps around the technology in the absence of additional legislation seeking to apply existing rules and legislation to the novel problems presented by AI. Given the potential for significant and life-altering consequences of its use in the financial sector, delays in such enforcement can have profound implications for some Americans.

It is promising, therefore, that the White House has released several executive orders on AI in addition to its AI Action Plan, which was published in July 2025. The administration is interested in securing its place in the AI race against global competitors. However, very little is mentioned to

²⁷ Finance, Competitiveness & Innovation Global Practice. “Global Experiences from Regulatory Sandboxes.” World Bank Group, 2020. <https://documents1.worldbank.org/curated/en/912001605241080935/pdf/Global-Experiences-from-Regulatory-Sandboxes.pdf>.

²⁸ Lee, Nicol Turner, Paul Resnick, and Genie Barton. “Algorithmic Bias Detection and Mitigation: Best Practices and Policies to Reduce Consumer Harms.” Brookings, May 22, 2019. <https://www.brookings.edu/articles/algorithmic-bias-detection-and-mitigation-best-practices-and-policies-to-reduce-consumer-harms/>.

address AI's use in financial services, where it is already being used by major players, like the Fair Isaac Corporation (FICO), that are involved in many consumer credit decisions. The CFPB is the federal regulator of FICO, but following the current administration's attempts to shutter the agency, important questions remain as to how algorithms will determine these scores.

Similarly, the Plan does little to address the discriminatory implications of these models that may reflect the historical lack of opportunity granted to racial minorities. As Brookings scholar Aaron Klein has shared in response to the current guidance, AI's use in mortgage underwriting might contribute to lack of growth in Black homeownership.²⁹ Yet, there is very little discussion on the underlying issues of representation and fairness in both training data and deployment.³⁰

Similarly, the rights of states may be compromised by the objectives of the AI Action Plan if state leaders exercise consumer protections for their constituents that federal officials appear to find burdensome. Earlier this summer, Congress considered, and ultimately struck down, a 10-year moratorium on states legislating AI, which threatened both states' rights and public interest.³¹ A bipartisan coalition voted 99-1 to strike it down, indicating support for allowing states to regulate AI, but the current AI Action Plan still suggests that AI-related funds should not go to states with "burdensome AI regulations" that are "unduly restrictive to innovation."³² Both Congress and states can exercise jurisdiction over consumer protection, and the latter, which is closer to consumers and

²⁹ Klein, Aaron. "Reducing Bias in AI-Based Financial Services." Brookings, July 10, 2020.

<https://www.brookings.edu/articles/reducing-bias-in-ai-based-financial-services/>.

³⁰ Friedler, Sorelle, Cameron F. Kerry, Aaron Klein, Raj Korpan, Ivan Lopez, Mark Muro, Chinasa T. Okolo, et al.

"What to Make of the Trump Administration's AI Action Plan." Brookings, July 31, 2025.

<https://www.brookings.edu/articles/what-to-make-of-the-trump-administrations-ai-action-plan/>.

³¹ Morgan, David, and David Shepardson. "US Senate Strikes AI Regulation Ban from Trump Megabill." *Reuters*, July 1,

2025. <https://www.reuters.com/legal/government/us-senate-strikes-ai-regulation-ban-trump-megabill-2025-07-01/>.

³² Trump, Donald J. "America's AI Action Plan." The White House Office of Science and Technology Policy, July 2025.

<https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

their local banking industries, should have some buffer from strict scrutiny over their desires to protect everyday people in this already highly regulated sector.

V. Proposed actions to make AI safe in financial services

As the use of AI grows in financial services, so should the oversight and guidance around its use. Existing and potentially new regulation can ensure that the industry's use of the technology is fair for consumers now and as the technology advances, allowing for more autonomous use cases. There are a few key steps that should be prioritized to address concerns not only in financial services, but across sectors with similar integrations of the technology that affect the sector.

1. Ensure the industry is compliant with existing legal statutes and remedies

The financial sector is already regulated by numerous federal and state organizations. It is imperative that industry be compliant with existing legal statutes and remedies that seek to prevent many of the problems and discriminatory results that have been present and could be furthered by AI. AI is a powerful tool but is not exempt from the rules and laws that govern our country. They must comply with transparency and nondiscrimination laws, particularly with consumer-facing AI. This means not just empowering agencies to carry out such enforcement but also providing them with the necessary funds to do so, perhaps reconsidering the role of the CFPB or other independent agencies in harm mitigation and increasing their budget for appropriate oversight. After experimenting with regulatory sandboxes in 2024, the CFPB determined that these initiatives ended up cementing advantages for a single market participant. To proceed with any proposals that combine financial services with AI experimentation, Congress must foster pro-consumer competition through an even playing field, not by handing out "waivers" that allow some firms to

skirt regulatory compliance.³³ Further, consumers are already guaranteed fair, transparent, and reliable decisions by various statutes, but they need to be enforced by independent agencies, and in some instances, state attorneys general. They must be empowered with the necessary resources to protect consumers in the age of AI, to maintain trust in our financial system, foster economic opportunity, and support individuals on their journeys toward financial stability. That also means that Congress must maintain the independence of federal agency and state decisions, and perhaps more on the latter in the absence of national leadership.

2. Mandate transparency guidelines and full disclosure for consumers

Companies that use AI in financial services should also be required to provide public consumer disclosures around AI use. These disclosure requirements should apply to firms using any algorithms or models that make decisions that can lead to adverse outcomes, even when humans make the final decision. Illinois has already amended a non-discrimination law to require employers provide notice to employees and applicants about AI's use in hiring decisions.³⁴ A similar law could apply to consumer finance models. Beyond transparency disclosures, model outputs should be explainable to accurately and reliably share with consumers how they came to a decision, pursuant to existing law.

Currently, researchers and labs are developing approaches to increase the reliability, reproducibility, and transparency of models, with differing degrees of success. Many frontier labs have been working on enhancing model reasoning, robustness, and interpretability. However, many

³³ Consumer Financial Protection Bureau. "CFPB Comment on Request for Information on Uses, Opportunities, and Risks of Artificial Intelligence in the Financial Services Sector," August 12, 2024.

<https://www.consumerfinance.gov/about-us/newsroom/cfpb-comment-on-request-for-information-on-uses-opportunities-and-risks-of-artificial-intelligence-in-the-financial-services-sector/>.

³⁴ Colvin, Jennifer L., Margaret R. Foss, Mallory Stumpf Zoia, and Samuel W. Newman. "New Illinois Labor and Employment-Related Laws Cover E-Verify, 'Captive Audience Meetings,' Noncompetition, AI, and More." Ogletree, April 14, 2025. <https://ogletree.com/insights-resources/blog-posts/new-illinois-labor-and-employment-related-laws-cover-e-verify-captive-audience-meetings-noncompetition-ai-and-more/>.

models, including those used in consumer finance, are still complex, black-box, and sometimes inconsistent, algorithms. If firms and developers do not fully understand how algorithms came to their decisions, consumers stand to suffer when they receive adverse decisions without explanation. Consumers deserve transparency as to when AI is being used, how it might be used to make eligibility decisions, or how it is involved in profiling one's economic readiness. These mandates would help provide agency to consumers who face adverse decisions and accelerate long-term adoption and trust in consumer finance models.

3. Encourage the responsible and ethical development of financial models

Since financial models that rely on AI are trained on traditional data sources that reflect existing societal disparities, biased models may impose disproportionate burdens on already vulnerable borrowers as already highlighted. In fact, almost 50 million U.S. adults lack enough credit history to be scored under common models. Black and Hispanic consumers are twice and 1.5 times as likely, respectively, to be unscored or considered to be subprime compared to white consumers.³⁵ Some researchers are investigating approaches to mitigate such bias and improve model accuracy, including diversifying data sources used to train models.³⁶ Model developers should collaborate with civil society, research institutions, and standards development bodies to adopt responsible AI development pipelines throughout the AI lifecycle, from using representative datasets during training to leveraging tools such as NIST's AI Risk Management Framework for proper assessments and technical standards for AI-enabled, financial products. Prior to deployment, model outputs should be validated for fairness, ethics, and transparency concerns. Further, financial organizations

³⁵ Koide, Melissa. "AI for Good: Research Insights from Financial Services." Center on Regulation and Markets at Brookings, August 2022. <https://www.brookings.edu/wp-content/uploads/2022/08/AI-for-good-Research-insights-from-financial-services-3.pdf>.

³⁶ FinRegLab. "Advancing the Credit Ecosystem: Machine Learning & Cash Flow Data in Consumer Underwriting," July 2025. <https://finreglab.org/research/advancing-the-credit-ecosystem-machine-learning-cash-flow-data-in-consumer-underwriting/>.

should be willing to engage in regulatory sandboxes that address anti-bias with some level of enforcement and accountability.

4. Brace for the adoption of agentic AI in the sector

As AI models continue to advance, newer, more autonomous models will enter the sector. Agentic AI systems, for example, are designed to operate more autonomously than traditional AI systems, with AI agents engaging with their environments to process information and complete complex tasks with minimal direction. Multi-agent systems (MAS) consist of multiple agents with different roles in the larger ecosystem: for example, one agent may retrieve data while another processes it or even coordinates other agents.³⁷ In financial services, agentic AI systems can be used for cybersecurity, compliance, and investment strategy tasks, as well as underwriting-related and claims adjustment purposes.³⁸ Agentic AI, which is poised to assist with more complex tasks, must be cautiously integrated into financial services to ensure the safety and security of consumer assets, and market viability. Since agentic AI introduced enhanced risks from increased autonomy, firms must retain control, oversight, and transparency over agents' decisions. Some best practices include continuous monitoring, mandatory expert-in-the-loop oversight, and legibility of model outputs.³⁹ Agentic AI must also only be used in appropriate contexts and model developers should place operational limitations on agents on the model-level, to ensure that default agent behaviors align with safety and risk mitigation principles.⁴⁰ Furthermore, frontier labs, standards development organizations, and research institutes must extend existing risk management and audit frameworks

³⁷ Quick, Jeanette, Colin Colter, Luke Dillingham, and Kelly Thompson Cochran. "The next Wave Arrives: Agentic AI in Financial Services." FinRegLab, September 2025. https://finreglab.org/wp-content/uploads/2025/09/FinRegLab_09-04-2025_The-Next-Wave-Arrives-Main.pdf.

³⁸ Ibid.

³⁹ Kwon, Joe. "AI Agents: Governing Autonomy in the Digital Age." Center for AI Policy, May 22, 2025. <https://www.centeraipolicy.org/work/ai-agents-governing-autonomy-in-the-digital-age>.

⁴⁰ Shavit, Yonadav, Sandhini Agarwal, Miles Brundage, Steven Adler, Cullen O'Keefe, Rosie Campbell, Teddy Lee, et al. "Practices for Governing Agentic AI Systems." OpenAI, December 14, 2023. <https://openai.com/index/practices-for-governing-agentic-ai-systems/>.

to ensure agentic AI frameworks are consistent with law, especially for public-facing applications. When models are fully autonomous in their decision-making, new forms of liability may arise, or existing approaches triggered. Without due diligence now, Congress may lose sight or fall behind in the compliance structure for fully autonomous financial decisions, which can start in financial counseling. Who makes the decision is, at best, a starting question for legislators, regulators, and industry representatives who are looking to embrace such sophistication in the sector.

5. Invest in AI financial literacy programs

Consumers must be provided with the means to understand the benefits, limits, and foundation of the technology. AI literacy must be widely available to increase agency and awareness among consumers in response to the accelerated offerings of AI in financial services. For all Americans to reap the benefits of AI in consumer finance, which range from providing faster customer service to expanding the accessibility of products at a lower cost, consumers must understand how to responsibly use and interpret AI outputs. AI literacy must also be paired with digital access and basic literacy on computers and the internet. Financial service activities can trigger consumer confusion when AI is intertwined with products and services. For individuals with limited access to the internet, or insufficient basic digital literacy training, they will be less likely to discern what is allowable to share in terms of their personal information, or how their data can be compromised if inputted into more public computer access terminals. With a persistent digital divide, it is important to accelerate more home broadband options for individuals to conduct private transactions from trusted spaces and internet-enabled devices, and to ensure that they understand how their data is accessed by financial institutions, and third parties. The recent elimination of the Digital Equity Act, which would have provided some foundational training in digital literacy sets society back when it comes to fluency in technology products, especially for rural residents, older

people, and those with disabilities.⁴¹ When AI is added in, these functions become more complicated to understand, and without national digital and AI literacy programs, more consumers will be victimized by scams and other fraudulent behaviors ranging from compromised identities to the loss of wealth, income, and necessary security assets. That is why Congress must consider how to support national initiatives aimed at advancing digital and AI fluencies, while advancing laws that protect consumer privacy, and ensure compliance of AI-enabled products and services with existing laws.

VI. Conclusion

We are at a critical point in providing protection to consumers facing AI's use in the financial sector. Whatever direction this takes will determine not just the next steps of innovation, but also the scale at which deployment and adoption are possible given the trust people will have in the technology. Though AI isn't new, there are still many unknowns. From the environmental impacts that AI can have on communities that are already strained to maintain clean air and water to the ramifications around labor and the effects of automated banking on some of the most vulnerable workers, there are still more questions than responses to AI's adoption and use in the financial sector. Providing consumers agency over how AI will be used to score them and determine their eligibility for an opportunity stands to benefit both the companies, who will better understand their consumers and their needs, as well as the public, who may grow more comfortable with the technology and lead to further adoption and efficiencies.

Thank you again to the Members of the Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence for the opportunity to testify. I also want to thank Brookings

⁴¹ Garner, Drew, and Grace Tepper. "The Digital Equity Act: What It Is and Why We Need It." Benton Institute for Broadband & Society, May 14, 2025. <https://www.benton.org/blog/digital-equity-act-what-it-and-why-we-need-it>.

researchers Josie Stewart and Michelle Du for their assistance in preparing my statement. I look forward to your questions.

Chairman STEIL. Thank you very much, Dr. Turner Lee.

We will now turn to member questions. I will recognize myself for 5 minutes for questioning.

I want to start with you, Dr. Cox, if I can. Let us do a real quick stage setting. Algorithms have been utilized in the financial services space since the 1980s, but we have now seen an explosive investment into AI. What is sparking this massive investment, in a short answer?

Mr. COX. Yes. I think what is sparking this massive excitement and investment is that this technology is general purpose in a way that previous technologies have not been. So there is a lot of work to train a system for, and now the same system can be used in many different settings.

Chairman STEIL. And it can pull through massive datasets in a way that was not possible in the 1980s.

Mr. COX. Absolutely.

Chairman STEIL. Although in the 1980s as well as today, there is already a regulatory framework in place, fair?

Mr. COX. Yes.

Chairman STEIL. Okay. Let me continue this. I want to come over to you, Mr. Gorfine, if I can. We are faced with something new as we focus on kind of risks and threats. The human mind is really good at thinking about what the risks are. It is not as good at thinking about the potential upside.

If we think about what the Biden Administration did, they really took an attitude to the extreme in its handling of artificial intelligence. Fortunately, President Trump has reversed course on this, issuing Executive Order 14179, Removing Barriers to American Leadership in AI, and releasing America's AI Action Plan in July 2025.

So this approach can really accelerate AI development through a try-first approach that removes some of the red tape and onerous regulations.

I want to come to you, Mr. Gorfine. In your opinion, how can the AI Action Plan—or how is the AI Action Plan an improvement over the previous administration's AI approach?

Mr. GORFINE. Thank you. Thank you for the question.

I think that what we are seeing right now is an effort to look where there is existing regulation in place and taking an approach of letting us allow this to develop so we can actually identify risks and determine whether anything more is needed.

I think, as many of us have outlined, the financial services industry is a heavily regulated industry. We have existing laws, regulations, and guidance that have been able to adapt and incorporate emerging technologies for decades. There are some approaches, including around the world, that are looking to do things kind of preemptively and in a more prescriptive fashion verse an approach where we say, hey, we have this scaffolding in place, we have principles in place to guide adoption of technologies in a responsible way, let us allow this to actually grow and develop.

I think it is critically important that the U.S. does remain the global leader in AI, including in the financial services context. So the environment seems set to unleash that.

Chairman STEIL. Thank you very much.

Mr. Reisman, I want to build on what Mr. Gorfine just said. In particular, let us dive into data privacy laws, something that is heavily regulated in the financial services space, but with the advent of AI, there are additional concerns that are coming online.

So as financial services firms are increasingly deploying AI in areas like lending, fraud detection, customer services, what role will Americans' personal financial data play, and how do we use AI to implicate existing data privacy laws? Of course, the logical follow up would be, what should we be looking at from an AI perspective as it relates to data privacy?

Mr. REISMAN. Thank you very much for this important question.

It has been heartening to me to hear how much we have talked about data privacy today, because I think what we would say is that a sound data privacy framework is a very important foundation for the development of AI and for America's AI leadership. I think data is the raw material on which artificial intelligence depends, and AI model developers depend on having rich datasets to develop high-quality models.

One thing I think that often gets put into the discussion is the idea of a binary between privacy or innovation, and we would absolutely say it is privacy and innovation, and privacy is an enabler. So a good data privacy law with modern interpretation holds onto classic principles we have had for decades in privacy law but allows for flexible use of data for purposes such as model training and development.

Chairman STEIL. So building on what you are saying, do you believe that we can regulate and address AI under the current regulatory framework or do you need a new regulatory framework, which has been proposed by some?

Mr. REISMAN. What I would say is that, as a first step, we should look at the framework that we have, look at it closely—and a lot of the elements, as has already been discussed today, are there—and try and be very precise about identifying any lacuna but first, let us look at what we have.

Chairman STEIL. Dr. Cox, you are nodding. You agree that we can in many ways regulate AI through the existing framework?

Mr. COX. I think, certainly, the starting place should be what are the risks, what are the outcomes that you are trying to govern and control, and then any modifications that come for the technology can come from there but starting fresh.

Chairman STEIL. Thank you very much. I agree.

I yield back.

I will now recognize the gentleman from Massachusetts, Mr. Lynch, who is also the ranking member of this subcommittee. You are now recognized for 5 minutes.

Mr. LYNCH. Thank you, Mr. Chairman.

This morning, a number of us had a meeting with Jack Clark, the co-founder and head of Policy at Anthropic and it was, I think, enlightening to hear him say how, for Anthropic, that the investment, the advancements, the deployment of AI had exceeded by about 5 years their expectations of the development and advance of AI, even based on the projections that they made 3 years ago. They said, we are 5 years beyond where we thought we would be today.

I just point to the velocity of change here in this sector and what a problem it creates up here, in terms of trying to regulate that and protect markets as well as protecting consumers.

I know a couple of you have mentioned the sandbox model. Ms. Turner Lee, the Singapore model—and we actually had a congressional delegation from this committee go to Singapore and kind of review the model that they had for their sandbox back when they were doing fintech, a fintech sandbox and a crypto sandbox.

So, interestingly enough, their sandbox has an ethical framework which begins with a principles-based approach so that there are elements of fairness and accountability, inclusivity, and it sets an ethical foundation for all AI projects. It has a bias mitigation element to their sandbox and mitigates bias in those AI systems, and developers are encouraged to assess their algorithms for potential discrimination.

There are data transparency practices that must be complied with. Organizations must be clear about how data is collected, processed, and shared—explainability requirements within their sandbox model. So they must ensure that their AI systems can provide understandable justifications for their decisionmaking.

There is also very strong stakeholder engagement, encouraging a wide range of stakeholders, and massive user privacy protections that are very important not only for Singapore but globally and certainly for the United States as well.

There are regular reviews and audits of that sandbox process—sandbox process among these participants in the sandbox, as well as blueprints and guidelines and knowledge sharing.

Is that the formula that we should use if we are going to go—I know the full committee chairman has an idea about sandboxes and we are going back and forth, he and I, about, what that should look like.

I would just like to get your feedback. You have touched on a lot of the sensitive issues, and I would like to hear you extrapolate a little bit more.

Ms. TURNER LEE. Thank you so much for that question.

I mean, I am a fan of sandboxes. I wrote a paper in 2018, when we were just beginning this bubble, on what algorithms were going to do when it came to consumers and regulatory sandboxes are a great way to experiment. In fact, here in the United States, we have used sandboxes to cultivate the fintech marketplace but because of the velocity of speed in which AI is gathering our personal information, the velocity in which we are going into transformational models where we cannot determine the beginning and the end. When we actually deploy a sandbox model in this day and age, it is important to have many of those verbals that the Singapore Government has put into place.

It allows for accountability, continuous monitoring, transparency, and it also allows us to ensure that consumers are baked into the process, as opposed to providing waivers and exceptions to companies to experiment and then come back and tell us how it is all going to work, which is the current model from which we exercise right now when it comes to AI.

Mr. LYNCH. Thank you. One problem that I noticed in the Singapore model was the number of participants was limited, so it never

scaled up. Some of the problems that we see with AI is when you get that scalability.

Is there a way for us to—I know the auditing and review after the sandbox is ongoing, but is there any way we can sort of make sure that scalability effect is implemented inside of a sandbox?

Ms. TURNER LEE. Well, I think one of the areas in which many of us who have fought for consumer protections are interested in—

Chairman STEIL. The gentleman's time is—I will ask you to write that for the record—

Ms. TURNER LEE. No problem.

Chairman STEIL [continuing]. cognizant of the time.

The gentleman yields, and the gentleman from Michigan, Mr. Huizenga, who is also the vice chair of the full committee, is recognized for 5 minutes.

Mr. HUIZENGA. Thank you, Chairman Steil.

Dr. Lau, we are talking about sandboxes. Do you have anything? I know you have been commenting on this in the past. Do you have any additional thoughts that you would like to—

Mr. LAU. Definitely. We are actually participants in the Singapore AI sandbox that they have set up, the AI Verify Foundation. It has actually been working exceedingly well, in the sense that these types of government agencies need to bring the latest technologies to actually evaluate the latest risks that are emerging with AI agents and new AI technologies.

So by bringing in innovative technologies to do that type of red teaming and testing, they are able to stay ahead. Think about here in the U.S., where we have regulatory agencies that are still being educated about the technology, whereas in Singapore they are able to unleash a lot more in terms of AI potential through bringing in new technology to test and evaluate it. I think there are a lot of benefits there and also to the point that Congressman Lynch mentioned around scalability, red teaming is a very labor-intensive problem, right. So, you have thousands of tests that enterprises have to perform. If the Federal agencies can provide more guidance and tooling around that, it is going to accelerate AI adoption.

Mr. HUIZENGA. Okay. That is helpful. Thank you.

Staying with you here, we often hear of an AI arms race with China that, if lost, would threaten U.S. national security and the United States global economic dominance.

In your opinion, is the United States currently—how are we currently faring, I guess, in this arms race? Are we failing at it? If so, how? More importantly, how can Congress and the U.S. Government writ large ensure that the United States outpaces China in the AI space?

Mr. LAU. I think there have been different paradigms emerging here. One is where you have a Federal Government that can choose and pick and invest in winners of the space versus creating an open marketplace where different, say, large language model providers or vendors can compete within this marketplace, right.

I think what we are seeing when we talk to folks within the Department of Defense (DOD), or Department of War now, we see actually a movement toward opening up these marketplaces where

many different vendors can compete and the best solution rises to the top.

Again, I will move back. We are actually evaluating which are the best solutions for use cases that is going to drive America's national security interests. It really comes down to can we do the right evaluation of those large language model providers or vendors.

So combining open marketplaces that foster innovation allow different folks to compete, and then we get to choose the best ones for our use cases. It is absolutely essential here for us to be competitive in this arms race you are talking about.

Mr. HUIZENGA. By the way, I deeply appreciate how you are succinctly wrapping up some very complicated and detailed things and allowing me 2 minutes yet to ask other questions. That is important up here when we only get 5 minutes.

Mr. GORFINE. I want to talk a little bit about the impact of U.S. regulation and how it may threaten the use of AI in the financial services space. We had the Securities and Exchange Commission's (SEC's) predictive data analytics proposal under the Gensler SEC, which I believe would have had some significant impact.

Are there other things like that we need to make sure that we are aware of and avoiding?

Mr. GORFINE. Yes. I appreciate you raising the predictive data analytics rule. I mean, that was rescinded and I think that was an example of a non—

Mr. HUIZENGA. Rightly so. Rightly so. As we saw with that SEC chair, he was expansive in his pursuit of covering every territory he possibly could.

Mr. GORFINE. Right. I would agree it was not technology-neutral, and that should be kind of a guiding principle. It had incredibly broad—

Mr. HUIZENGA. So it should be tech-neutral.

Mr. GORFINE. Absolutely should be technology-neutral and importantly, like the soft power too and the way that regulators communicate to the marketplace is important. So what the—

Mr. HUIZENGA. Oh, in other words, you do not want to be pounded into the ground and threatened with being put out of business? That does not foster innovation?

Mr. GORFINE. That can be a challenge.

Mr. HUIZENGA. Allow the sarcasm to be mine.

Mr. GORFINE. That can certainly be a challenging environment, I think, especially when you are talking about smaller firms. When I think about regulatory messaging, it is the small firms, it is the community banks that do not have large armies of compliance teams that are able to parse certain types of messaging that comes from regulators. So it can serve as a big deterrent to adoption.

Mr. HUIZENGA. In the last 30 seconds, how can these smaller community financial—you mentioned community banks or credit unions. How can they use AI to compete?

Mr. GORFINE. Well, so I am a big believer—and I say this in my written testimony—in standard setting organizations. Like, industry-driven standards that regulators can effectively recognize or even create safe hardware would allow small firms to know that, if you are compliant with these best practices or standards, that is,

from a diligence perspective, a route you can pursue. I think that exploring those types of models in the U.S. is going to be critically important as—I am out of time.

Mr. HUIZENGA. I yield back. Thank you.

Chairman STEIL. The gentleman yields back.

The gentleman from Illinois, Dr. Foster, the ranking member of the Financial Institutions Subcommittee, is recognized for 5 minutes.

Mr. FOSTER. Thank you, Mr. Chairman, and to our witnesses.

I guess Mr. Reisman and maybe others have mentioned the importance of privacy enhancement techniques. Several years ago when I was chairing the AI Task Force on this committee, we had dragged in a witness to talk about homomorphic encryption and some of the differential privacy techniques and I noticed just recently—and it was clear back then that these were not ready for prime time. There was a huge penalty for performance, and the privacy was not actually that great.

I noticed recently Google just announced this Google Vault Gemini, which sounds like they are actually implementing training with differential privacy.

So I was wondering if anyone on the committee could say something about are these techniques really ready for prime time, and is there the possibility of a regulatory safe harbor for firms that commit to using high-quality differential privacy type tools?

Mr. REISMAN. Thank you for the question.

It has been quite extraordinary to see the progress on privacy enhancing technologies over just the last few years. I think you are right that not so long ago many of them were not ready for prime time, but the curve has been quite steep in terms of many of them becoming much more feasible.

Part of it is that many privacy enhancing technologies are compute-intensive, but our computing power collectively is growing a lot stronger, so we are able to handle that. I think there was also a sense that a lot of them were complex and maybe out of reach, especially for smaller companies.

There is a whole ecosystem now of expert companies that are able to consult and offer what you might call off-the-shelf privacy enhancing technologies, to make them much more available much more broadly. So—

Mr. FOSTER. It would make sandboxes easier to implement.

Any other sort of comments on the state-of-the-art? Would anyone disagree that this is something that is pretty promising at this point?

Mr. COX. Things like fully homomorphic encryption provide very strong guarantees and you are correct that they are slower than conventional computing, but the leaps and bounds that they are making in that technology, it is orders of magnitude, thousand times better. Every time I turn around, it has gotten better.

So I think that is moving very fast, and the harbor is going to ultimately start to close that gap as well. So I think we are going to benefit from that wave.

Mr. FOSTER. That could be a key component to a safe sandbox in a lot of these things.

Mr. COX. Absolutely.

Mr. FOSTER. Relating to one of the things we struggle with is the small bank versus large bank, trying to level the playing field if we can and it is one of the things that AI can potentially help us with.

Soon all of us will have in our pocket the best team of lawyers that has ever been assembled. So it means if you ever get in a fight with a billionaire, a legal fight, the billionaire will not be able to get a better legal team than you have essentially for free.

Now, similarly, a small bank should, through AI, be able to have access to a really good AI risk adviser, a team of risk adviser AIs, so that even if it is a pretty small bank, this team will know every way that any bank has failed in the history of this country and just have a tremendous amount of knowledge and make it easier to operate a small bank.

Similarly, regulations, regulatory technology (RegTech), I think, is another real advantage. I did an interesting experiment a couple weeks ago where I said, okay, Claude, or whoever I was using, give me the balance sheets for three banks that have just failed for typical reasons, and it just knocked it out of the park.

Then I said, now give us a resolution plan for each of those three and it was great. I mean, it just said, okay, in this case set up a bridge bank. In this case, try to merge it. In this case, just liquidate it. They were very sophisticated things, and I was impressed.

I was wondering if any of you have a feeling for whether that is really kind of the future of regulation, that even small banks have all of their records electronic.

If there was a standard interface for the accounting software that gets run that could report up to risk management and to the regulators, you could have real-time stress testing against dozens of scenarios every single night for the smallest bank, and that would be a real leveler.

Any thoughts? Is there a reason why things cannot evolve in that direction?

Dr. LAU.

Mr. LAU. Thank you. Yes, certainly. We work with a lot of community banks and regional banks, backed by a number of them and we see this every day, right. So this new AI technology can empower them to automate or streamline a lot of these compliance work flows that are very manually intensive that they just do not have staffing to do, but also, as you mentioned, unlock really new types of auditing and risk controls, so 24-hour continuous monitoring and observability into certain types of work flows that are highly regulated. That is all really exciting and actually being realized/materialized today across a lot of these smaller banks but I would say on the flip side, you also want to make sure that as these AI technologies enter into these regulated work flows, they are also having the right guardrails in place and they are being used for those work flows appropriately and aligning them with proper bank policies and procedures. So that is something that actually requires technical expertise.

Chairman STEIL. The gentleman's time is expired. We ask you to complete that in writing.

The gentleman from Ohio, Mr. Davidson, who is also the chair of the Subcommittee on National Security, Illicit Finance, and International Financial Institutions, is recognized for 5 minutes.

Mr. DAVIDSON. Thank you, Chairman.

Clearly, AI represents a transformative force, one that can exponentially enhance productivity, bolster our financial system's global edge, and it will no doubt influence our culture. Technology never exists in a vacuum, and I have long warned about the erosion of our Fourth Amendment protections in the digital age.

The Uniting and Strengthening America by Providing Appropriate Tools Required to Intercept and Obstruct Terrorism (PATRIOT) Act, for example, massively expanded domestic surveillance. The Bank Secrecy Act had long ago obliterated any real claim to privacy in your financial dealings. Most State bureaus of motor vehicles are monetizing the personal data citizens are required to provide just to drive or obtain ID. Now, the government is outright buying data that would otherwise require a warrant or a subpoena, sidestepping the Fourth Amendment entirely.

Because we are serious about fostering innovation, and we are serious about our Constitution, we also need to recognize that AI should serve the American people without turning it into another tool for unchecked surveillance or data exploitation.

Mr. Chairman, I would like to submit this document for the record. In my recent op-ed for the Daily Caller, I emphasize privacy is the foundational layer for ethical AI.

Chairman STEIL. Without objection.

[The information referred to can be found in the appendix on page 96.]

Mr. DAVIDSON. Thank you, Chairman.

Without it, we are handing the keys over to a surveillance state on steroids. I mean, we need to update the regulatory framework for AI or build it in Congress, and we certainly need to address privacy, because that is the dataset that AI is using. Urgent questions, including who is liable if AI is misused? Who profits and how when someone else's data or intellectual property is indexed or shared with AI? When does law enforcement need a warrant, if ever?

These are not abstract questions. They are urgent, and Congress needs to be proactive, not reactive, to set clear boundaries. We cannot let AI become another excuse for Big Brother to pry into our lives.

So what safeguards do we need? We should learn from other jurisdictions, but we should, of course, chart our own course. The EU's AI Act, for instance, takes an approach with outright bans on real-time biometric surveillance and social scoring, measures that align with protecting individual liberties from invasive tech but in other ways, the European Union has become pretty Orwellian with some of their privacy approaches and speech limitations.

So positively framing safeguards around AI is important. Freedom surrendered is rarely reclaimed, so I am encouraged by the Trump Administration's AI Action Plan with its emphasis on accelerating innovation, building infrastructure, and leading globally. It promotes open source development, cuts red tape, and streamlines permitting without smothering the private sector.

Mr. Reisman, how are financial regulators protecting the vast amounts of sensitive financial data that the government already collects from AI exploitation?

Mr. REISMAN. Thank you for the question.

First of all, I just want to—to your point about raising data privacy and this very important foundation for a sound environment, not only for keeping people safe but for innovation. I would underscore our sense that having a sound privacy framework in the United States is actually not only consistent with but important for achieving the goals of the AI Action Plan. It is encouraging to see progress in that area.

I think others may have more experience than me with what is going on inside the agencies for what they are doing to keep data safe, so I would defer to others.

Mr. DAVIDSON. Anyone else?

Mr. GORFINE. I am happy to take that. I would agree that data is the one place screaming for Federal legislation to create a proper baseline for how we secure data that will be consumed by AI models.

With respect to your question, what can regulators do? We were just talking about privacy enhancing technologies, things like encrypting data. Ultimately, financial regulators need to be recognizing that there are these developing privacy enhancing tools and making sure that regulated financial institutions can use those technologies.

So data can be encrypted. We can avoid creating honey pots of information that are being sent every which way. I mean, the reality is today all of our information—driver's license photos have been emailed, sent to every single provider. There is a better way to move encrypted information and limit access to who actually gets it and when.

Mr. DAVIDSON. Thank you. I think the computing architecture is really important. I have always liked blockchain technologies. I like zero knowledge proofs and ways to protect data that way. The government in some cases has artificially limited that.

I will submit written questions for the record. I am really curious how you protect intellectual property, copyrights, trademarks, and maybe monetize that for other people down the way. Maybe blockchain does it. Maybe there are other ways but thank you for your expertise and attention.

Thanks for this hearing, Chairman, and I yield back.

Chairman STEIL. The gentleman yields back.

The gentlewoman from California, the ranking member of the full committee, Ms. Waters, is now recognized for 5 minutes.

Ms. WATERS. Well, thank you very much.

Dr. Turner Lee, we have seen development of new AI technologies like agentic AI, which can operate independently and autonomously. Agentic AI is capable of self-directed and complex behavior and can also carry out multi-step tasks.

Financial services companies are now experiencing and experimenting with agentic AI in investments, credit decisioning, and more. However, I am concerned about the new scale of risk and vulnerabilities from all of this.

Dr. Turner, can you elaborate on the potential risks for agentic AI and whether a liability framework can help us mitigate those risks? Who is responsible for the actions of an AI agent if something goes wrong?

Ms. TURNER LEE. Thank you so much, Congresswoman, for that question.

I think most of the conversation we have had today is on this promise of AI as it relates to efficiency in the financial sector and agentic AI as an extension of that. The more and more financial institutions can become more autonomous through chatbots and other tools that allow people to not necessarily talk to a financial counselor but to talk to AI to be able to get through their steps. It makes sense on the productivity side, but it does not make sense for the consumer.

The consumer has a lot of reputational risk that comes with this, a lot of financial risk and in an industry where we know that consumers are at the heart of any type of harm or potential harms that come through misadvice or miscalculation, it could have a huge effect not only on people who are wealthy but people who are experiencing the wealth gap.

So I appreciate your question. I do think we need a liability risk structure for this, one in which Congress thinks through the same type of consumer protections that we are defunding right now and diminishing, even with many of the actions today to append those regulatory agencies like the FTC and the CFPB.

As we move into agentic AI, without those guardrails and the ability of States, in particular, to control how people respond to this new technology, it is actually going to go further than we can catch up with it.

Ms. WATERS. Wow.

Dr. Turner Lee, the Republican regulatory sandbox proposal for AI we are considering today appears to be just more deregulation framed as innovation and lacks requirements for public disclosure, harm mitigation, and other important protections, with an unlimited scope and virtually no limitation, putting Americans, consumers, and other market participants at risk.

I am deeply concerned that regulatory sandboxes may remove safeguards from a rapidly developing AI market that is already lacking meaningful Federal regulations and oversight. In fact, we have yet to fully understand the consequences of deploying this technology without any safeguards and are already dealing with countless public safety issues.

Despite the fact that you have alluded to some of this, what risks do regulatory sandboxes pose? If Congress were to consider regulatory sandboxes for AI, what kind of standards should responsible sandboxes meet?

Ms. TURNER LEE. Thank you for that question as well.

I mean, I think we have mentioned already the role of sandboxes in helping us to cultivate new products and services within the sector—and among other sectors. We actually see sandboxes in healthcare.

The challenge is, without the right variables that we are actually constructing to make sure they are safe, ethical, fair, inclusive, as

has been mentioned, sandboxes will turn out to be an exceptional exploitation of consumers.

What that means is, without any guardrails—which, the AI Action Plan is suggesting there should be modifications, there should be waivers. We really need sandboxes to be very transparent. We need clear goals and questions about what it is trying to solve. We need protections against consumers who are part of those sandboxes to ensure that any harm that they face is—there is retribution for that as well.

In my opinion, we can actually create this without putting people at risk in terms of the end product, and we can do it in the light, as opposed to in the dark, when it comes to companies and government working together on those policy concerns.

Ms. WATERS. I am so pleased that you are here today.

Ms. TURNER LEE. Oh, thank you and I am pleased you are here today too.

Ms. WATERS. I am going to call on you to continue the kind of education that is needed with Members not only of this committee but this entire Congress. I am particularly concerned about the possibility for wrongdoing, the possibility for discrimination—

Ms. TURNER LEE. Yes.

Ms. WATERS [continuing]. all of that and I want to learn more about the databases that are being used and how they are significant in determining the outcomes.

Ms. TURNER LEE. Yes.

Ms. WATERS. Thank you so very much.

Ms. TURNER LEE. Thank you.

Chairman STEIL. The gentlewoman yields back.

The gentleman from Tennessee, Mr. Rose, is recognized for 5 minutes.

Mr. ROSE. Thank you, Chairman Steil, and thanks to Ranking Member Lynch for holding this important hearing.

Thank you to all of our witnesses for taking time to be with us today.

I want to start with you, Dr. Cox. Could you please describe the concept of technological singularity and share your perspective on whether you believe AI will achieve singularity and provide an estimate on the possible timeframe for this development?

Mr. COX. Thank you for the question.

So the idea of a singularity is, we reach a point where the technology is able to advance its own progress faster and faster and faster, such that it can get sort of a positive feedback loop and then we suddenly have an explosion of capability.

One of the things that is interesting that—I just dusted off a copy of “The Singularity is Near” by Ray Kurzweil that happened to be on a shelf that I was cleaning up. One thing you will see about futurists is, they often get the shape of things right, but the years are difficult to predict. I would not hazard to make a guess about when that is going to happen or if it is going to happen but I would just say that I do not think we are anywhere near—as somebody who works with this technology day-in and day-out, I do not think our biggest risks come from, sort of, some eclipse of “AI is better at everything than humans are,” but certainly our labor-market issues that we have to deal with as sub-tasks of a job are

displaced but I do not think we are in imminent danger of anything so extreme as a singularity.

Mr. ROSE. Thanks. I appreciate the insight.

Dr. Lee, I found an interesting article from the nonprofit Cash Essentials that states, quote, Some AI-driven systems may treat cash transactions as suspicious, leading to increased scrutiny and regulatory pressure that discourages cash use, unquote.

How can we ensure that AI does not discriminate against individuals who use cash? Additionally, what measures can be taken to prevent AI systems from falsely flagging cash transactions as suspicious, thereby avoiding undue pressure on companies and regulators to favor cashless payments over cash transactions?

Ms. TURNER LEE. I do appreciate that question because I think we are seeing, based on the regulatory sandbox that we used to create the fintech industry, a lot more transactions, especially among the unbanked or underbanked.

To your point, here is why I think AI could be an interesting tool to help us combat fraud. Improved analysis through AI detection systems could actually be helpful there. Using AI in ways that the AI sort of cleaves with the data about the underbanked or unbanked and combines that with traditional legacy systems could also be helpful.

We also talk about AI as just taking people's data, but it also considers other externalities, like where you live, what your ZIP Code is, these other proxies. Oftentimes those proxies are for the worse, in terms of bias and discrimination, but they could actually also be helpful in helping us understand what the new, modern economy looks like in banking.

So I would suggest that there are some techniques on the technology side, but there is also room for people like me, as a sociologist, to come sit at the table and help determine how we do this better.

Mr. ROSE. We already see—without the benefit of AI, we see cash transactions being discriminated against in a number of ways. It is a very real concern to me.

Mr. Reisman, with the increasing prevalence of artificial intelligence, it is clear that AI-assisted fraud will likely escalate rapidly, employing novel tactics to deceive consumers. In my view, it is essential to warn the public about emerging AI-driven fraud schemes as soon as they are identified.

How can private companies and industry stakeholders collaborate proactively to anticipate and combat the evolving threat of AI-assisted fraud?

Mr. REISMAN. Thank you for the question.

I share your concern about fraud, and I think I have heard a lot of folks say we are in a moment where the defense has to keep up with the offense, right? We want to make sure that our financial institutions are empowered to have the same tools to detect fraud that the fraudsters are using to advance it.

That is the most important, is that our institutions, whether they are small banks, whether they are large ones, have access to that state-of-the-art screening technology.

I think—one thing that I talked about in my opening testimony was the importance of dialog and I think making sure that we have

spaces, whether it is through sandboxes or whether it is through other mechanisms that are set up by the Congress or by the regulatory agencies, to constantly be sharing information about these advances in technologies and about these fraud capabilities so that we are all collectively working together on solutions.

Mr. ROSE. Thank you very much. I appreciate those insights. I yield back.

Chairman STEIL. The gentleman yields back.

The gentleman from California, Mr. Liccardo, is recognized for 5 minutes.

Mr. LICCARDI. Thank you, Mr. Chair.

Thank you all for your testimony and for taking the time to help educate us, myself in particular. I know this is a fast-moving area and I really had questions targeting really primarily Dr. Cox and Dr. Lee. I would be interested in your views.

We have several concepts that have been presented in legislation here before us, and I wanted to see if there was—I would like to tinker a little bit with a couple of these bills from Chairman Hill, the sandbox concept, and from my colleague Congresswoman Pettersen, the task force concept, and combine them in a way, if we could.

I represent Silicon Valley. Obviously, we have a lot of concerns—or I hear lots of concerns in my neck of the woods about having government involved in regulating the code and the algorithms, particularly since government is not terribly good at it. We would expect that the technology is evolving so quickly that we will not be good at it but, on the other hand, I think there is a widespread embrace of the notion that we need to be mitigating the worst of the harms.

So the question would be: If we were to create perhaps even more than a task force but an independent body of folks in industry, academics, experts, financial regulators, and others that were to establish what the best practices are in the industry—the best practices around, for example, curated data sets that I know you mentioned, Doctor, that IBM utilizes, or testing and reporting, or fraud detection, watermarking, data security—a whole host of best practices in the industry, and then hold that up as the standard, and establish, essentially, if that is going to be the negligence standard—or, essentially, the standard, if everyone complies with that standard, then you are exempt from liability. If you are not meeting that best-practices standard, then good luck with the lawyers and the regulators.

What about an approach that would combine those two?

I would ask either Dr. Lee or Dr. Cox, if you would like to jump in.

Ms. TURNER LEE. I will jump in first.

Mr. LICCARDI. Yes.

Ms. TURNER LEE. I think that is a tremendous idea that we should consider, because we seem to be at the same stalemate every time we talk about these issues.

There is a technical side of it, and then there is this consumer output, and I do think that government regulators are good at the latter, the consumer protection side of it, because, guess what, that

is the side where our constituents are coming to us and saying, "Something has harmed me."

With regards to what you are saying, I would just like to offer to you that we did start that process under the previous administration. The Blueprint for an AI Bill of Rights was actually a nice glide path for getting to the same protections that you are speaking of, as well as collaboration. The executive order, which was soon after appended, had a lot of that conversation on how do you actually bring different bodies together from various disciplines, industry sectors, government, and civil society so that we actually solve this together.

To date, we have seen a lot that sort of eroded in the new administration as well as in the AI Action Plan, where we are primarily competing against ourselves, when we actually just see China as our only force of nature.

So I would suggest—I agree with you. I think there has to be more conversation, more collaboration, more disclosure among the various entities to get to the space that you are talking about and I think a task force would not be a bad idea. That was also recommended under the prior administration.

Mr. LICCARDO. Thank you, Doctor.

Dr. Cox?

Mr. COX. I think that there are some interesting things emerging already. The ISO 42001 standard for just the entire creation process of a large language model is one example of something emerging, and I think that gives you a little bit of a sense of how these things are playing out. It is a voluntary standard, you get audited against it, and it is some sort of mark of whether you have the proper hygiene there.

Now, that is different than regulating the algorithm, though, right? Like, it is—

Mr. LICCARDO. Right.

Mr. COX [continuing]. more about regulating processes and controls—

Mr. LICCARDO. Right.

Mr. COX [continuing]. and documentation and auditing, as opposed to saying, this technology, this algorithm is intrinsically worrisome somehow.

I think that is one thing that I think many of the panelists have raised, is that there is nothing intrinsic about the technology; it is, how are you using it? It always has to be in the context of how it is being used and we have use-and-risk-based regulation. That is a normal thing that we already have. So combining that with things like cybersecurity hygiene standards, which are now then being transported to the world of AI, I think that is—I think that is an evolution of what we already have to a good outcome.

Mr. LICCARDO. All right. Thank you. Thank you both.

I yield.

Chairman STEIL. The gentleman yields back.

The gentleman from Montana, Mr. Downing, is recognized for 5 minutes.

Mr. DOWNING. Thank you, Mr. Chair.

Thank you to the witnesses.

This is a really exciting topic for me. The thoughts of AI—I came out of technology—the opportunities are incredibly exciting, but it is also a little bit scary in where it goes, and it is interesting—I was talking to Eric Schmidt last year. He wrote a book on AI with the late Henry Kissinger, and he made a comment that I thought was really interesting, about being very light on how you regulate it so that we are not blown out of the water by our competition. Then he said something that really caught my attention: “But you need to be ready to unplug it.” I was not exactly sure what that meant, but it was an interesting comment that I am still dwelling on.

As a former regulator, we had to deal with a lot of the issues with artificial intelligence as a tool for industry and really what the implications were. Some of the interesting things we did as an insurance regulator is—we were of the opinion that, once you wrote down a rule or a law or something prescriptive, it was probably already stale, because this is just evolving so quickly.

So what we tried to do to inform industry on how we were looking at it is, all the work we did on that committee was—it was concept space, non-prescriptive, because you wanted to kind of guide how we as regulators were looking at it without saying, “This is what you have to do.”

A lot of the questions that came up there were is the machine biased? Is the machine bad? Is the machine—a lot of this kind of concern about what was coming in it.

You know, me, you know, I think about the data. As a former researcher, I think garbage in, garbage out. What is the data you are training it with and what are the thought processes on how broad and how deep that data set is and what kind of results you have?

But the results coming out of it—a lot of folks were saying, “Well, at least in terms of a regulator, we need to have causation and not correlation.” You know, I do not agree. I think if you have a high statistical probability of getting the same results from this system—a lot of our life is based on correlation and not causation—I think you can use that as reality or close to reality.

Then, if it is giving you a result you do not like, like it is affecting a protected class, if it is something that you do not like, that is the public-policy decision, not “machine bad.” Then you have the real, honest conversation about that public policy.

Another thing that folks would say on the regulator side very often is, “Well, we need to be able to look into that black box.” I think about, well, that black box—once you have an N-equals-close-to-infinity access neural network, no human can understand that. What tools do you need to understand that? You know—this is my opinion—the tool you need to understand that is probably generated through AI.

So you have AI doing that, and it comes back to the ancient question of, “Quis custodiet ipsos custodes?” “Who is watching the watchers?” You know? It is an interesting problem to come up with.

I am going to start—sorry about—thank you for indulging me on that.

One of the things that I think about a lot is, as a State regulator, States’ rights are very important to me and there has been a big

conversation about that. I strongly favor allowing States to lead on regulation where feasible but at the same time, it is essential that this technology has the legal and regulatory flexibility to continue innovating and develop in the United States.

So I am going to start with Mr. Reisman here.

What role do you think the States should play in regulating AI or are you of the opinion that this is something that Congress needs to tackle?

Mr. REISMAN. Thank you for the question.

I think the first question that we should be asking is, do we need more regulation on AI at all? We have had an interesting discussion about looking at what powers we already have, whether that is under Federal law or under State law, that allows us to address a lot of the questions that we may have about AI.

I think that there is—one thing that we have seen is, for both consumers and for businesses, it can be hard when there is a kaleidoscope of different State regulations and so there is a value in having interoperable Federal standards. I think there is an interesting question to be explored further about whether there may be particular elements that States need to address as a—

Mr. DOWNING. Right.

Mr. REISMAN [continuing]. complement but not substitute for that.

Mr. DOWNING. Yes. Thank you on that.

In the interest of time, I am going to move on.

The Biden Administration seemed determined to stifle AI in any way it could, focusing more on potential threats than its clear benefits and I do believe the benefits are strong.

So I am going to move to Mr. Gorfine, please.

Could you describe some of the most concerning aspects of the Biden Administration's approach to AI, from executive orders to a so-called AI Bill of Rights, and discuss what the impact would have been had Congress passed what the Biden Administration proposed?

Chairman STEIL. The gentleman's time has expired, but we will ask the witness to provide that for the record.

Chairman STEIL. We thank you.

The gentleman yields back.

Mr. DOWNING. On that, I yield. Thank you, Chair.

Chairman STEIL. The gentlewoman from Massachusetts, Ms. Pressley, is recognized now for 5 minutes.

Ms. PRESSLEY. Thank you.

Thank you to our witnesses for joining us today for what is really a timely hearing.

When I am in my community in the Massachusetts Seventh, people are very animated about the future of artificial intelligence. I hear it all, ranging from fears of AI bias to excitement for innovative opportunities, to concerns about implications on the future of work, to confusion about what is next. I am proud that we are confronting these issues head-on today.

September 26 to October 3 marks Boston AI Week. Startups, investors, researchers, and students will all convene to explore the role they play in the evolving AI landscape. Leaders in AI, like the Mass Technology Leadership Council, will host educational and

networking events to build on State investments in the next generalization of leaders, educators, and workers in this rapidly growing field.

I am proud to represent a district that is trailblazing the development of the AI industry. In Massachusetts, jobs in the technology sector make up 14 percent of our labor force, compared to 10 percent nationally and we know diversity in AI jobs is essential to confront bias, to maximize opportunity, and to make good decisions that help everyone benefit from these cutting-edge technologies.

Dr. Turner Lee, I want to focus in on your research about the role diverse teams play in AI development and deployment. What are some ways diverse teams may be helpful to promote ethical and innovative uses of AI?

Ms. TURNER LEE. Thank you so much for that and I congratulate you on the AI Week in Boston.

So I would just say there that it is important to have representative groups for a couple of reasons that we have spoken about today. I think, on both sides of this conversation, there is this tremendous excitement for the opportunities and then a tidbit of fear—and, depending on who you are, more or less, right?—based on what AI can do.

It is important that we have people who have the lived experiences of the variety of impacts that AI can have. Why that is important, with any other internet technology that we have had, is because AI, particularly the financial sector, is dispersed to be very individual to that person.

So, when we talk about data, data that is traumatized or discriminatory, that actually takes the account of the wealth gap that is experienced by, for example, Black populations, where they have been denied credit, loans, and other eligibility, shows up in the data. When you do not have people who understand that lived experience, that experience passes on generationally, and the AI tends to then affect their quality of life. Loan denials continue, et cetera.

So I think it is really important to have not only background diversity but diversity of various people from various disciplines, various sectors, sitting at the table. How are we building financial-sector AI without people who actually understand how community banking works or sit in the roles of experienced financial literacy counselors, for example? You have to have everybody there.

Ms. PRESSLEY. Thank you.

Dr. Turner Lee, how do you believe the Federal Government has a role to play in this and what is that, to ensure that people from all walks of life—women, people of color, low-income folks—can have careers in AI, especially in the financial services sector?

Ms. TURNER LEE. Well, I think you pointed it out really well when you talked about trying to empower local entrepreneurs. It is important for the Federal Government to create low benchmarks to entry. We see on this panel a young man who has his own company that is doing incredible things when it comes to data architecture, AI infrastructure, et cetera.

Giving opportunities for a variety of people to participate in this ecosystem is really important, and I think the government could actually find ways to incentivize the private sector to be more in-

clusive of diverse founders and entrepreneurs and small businesses that want to participate in this space.

I also think the Federal Government could do more on AI literacy, what does a national AI literacy initiative look like, so that people know that this is not about being an engineer; this is about actually experiencing this behaviorally. The way that AI is dispersing, this is not going to be just about experiencing it on your computer; it is going to actually show up in your refrigerator and other places. People need to know that.

I would say the other thing that Congress should do is close the digital divide. We keep talking about AI, but we actually have not closed the basic infrastructure issue. We cannot build the compute facilities and data centers without that.

Ms. PRESSLEY. Thank you, Doctor, for being so prescriptive there.

Whether it is a hearing in Congress or a meet-up at the Museum of Science, AI is the topic of conversation, and we have the responsibility to ensure that all voices are included, because AI works best when it works for all.

I yield back.

Mr. DOWNING [presiding]. The gentlewoman yields.

The gentleman from Iowa, Mr. Nunn, is now recognized for 5 minutes.

Mr. NUNN. Well, thank you, Mr. Chair.

I want to thank the panel for being here.

Chair, as a fellow Air Force fellow, it is good to see you in that seat.

As we look at what Congress has been trying to grapple with, this idea of artificial intelligence—I have only been here two terms, and we have talked about it every single year. In fact, we were on the AI Task Force, something this committee helped lead to be able to address real solutions but I think Washington does a lot of talking, and we should be doing a lot more listening, particularly to practitioners in the field who have been addressing this—combating it, addressing it, and finding new solutions for it—and incorporating your experience into the way that we do policy. As opposed to DC. being the one who writes the rules of the road and then expects you to execute them, we should be doing this in collaboration.

Mr. Chair, it is one of the reasons I am leading a piece of policy called the Artificial Intelligence Practices, Logistics, Actions, and Necessities (AI PLAN) Act, or the Artificial Intelligence PLAN Act, to ensure that we here in Washington are incorporating and getting our own house in order before we start telling the innovators, the pioneers, and, candidly, the defenders how best to do it under a government auspice.

I look back at where we were with cybersecurity when I was Director of Cybersecurity at the National Security Council (NSC). Many of the solutions were already being provided out of the public sector and the private sector when we worked together.

With that, I want to be able to dive into a couple of issues that have been addressed when it comes to artificial intelligence.

First and foremost, the idea of misinformation and how it happens not only in the financial space but across our government. It has lowered the barrier to entry. It has allowed misinformation to

expand at an accelerated rate. We saw even this week, tragically, with the death of Charlie Kirk, Russian bots providing inception for everything from conspiracy theories to misaligning our own communications right here in our own country.

Dr. Christian Lau, you have been a leader at Dynamo AI. You have made an effort to really bring this technology on board. I want to ask you, what should Congress be doing right now to do those things like combating state-sponsored terrorism, including information warfare?

Mr. LAU. Yes. That is a great question.

Particularly, you mentioned the PLAN Act, which you proposed, which I think is a really great step forward in looking at these different risks—how AI can be used for misinformation, but also used for different types of financial crimes, right?—and threaten systemic risk to our system.

One thing that I would emphasize is that keeping up with the pace of the technology is key, and talking to practitioners about the latest risks is absolutely essential to making effective legislation, right?

One thing that I had an opportunity to talk to your team about is the advent of AI agents and particularly, you brought up cybersecurity risks—

Mr. NUNN. Right.

Mr. LAU [continuing]. right?

So AI agents is where you give these systems more autonomy to carry out end-to-end workflows and tasks. A lot of times, they can go in many different directions, but this also opens up new risks for agents to potentially exfiltrate data or for State actors to leverage agents and actually send what we call “prompt injection” or “jailbreak” attacks to actually exfiltrate information from protected systems in financial services or elsewhere.

So an AI agent can pull an email, that email could have an embedded hidden instruction or prompt injection, and that agent can then be overridden to actually perform malicious tasks like crash a system or exfiltrate sensitive information.

So I think the work that you have been doing on this act, as well as talking to leading practitioners, is absolutely key to keeping up with these risks.

Mr. NUNN. I could not agree with you more that the challenge here is, the technology is going to move much faster than the policy is able to keep up. So being able to provide a framework versus prescription is something the AI PLAN Act intends to do but also brings in best practitioners in this.

I want to talk with you, Matt Reisman. You are at the Center for Information Policy Leadership. You have seen some of these innovations come on board. Talk to us a little bit about how not only we incorporate them to have a better safeguard for our government and our citizenry but, where the innovation is really being pioneered a lot of times by our private-sector partners, how should we be incorporating that into how we work here in Washington, DC.

Mr. REISMAN. Thank you for the question.

I think, number one, one thing that has been encouraging to us to see is that, on the one hand—you have said, the technology is moving quickly. Responsible actors in this space are continuing to

innovate, but they are not leaving their values—and they are concerned about making trustworthy AI—behind.

Smart businesses in this sector, including a lot of the leaders, recognize that having trustworthy technology is not just good; it is good for business because customers, government, other stakeholders have more trust and faith in the technology that we are building—that they are building.

There is a lot to be said for government setting certain standards around transparency and asking for folks to show how does the technology work at a high level, and how are some of those specifications made.

Mr. NUNN. I would like to continue to see this level of collaboration, in the same way we have done with cyber defense and bringing an AI consortium together of business leaders like yourself, of innovators like yourself, to be able to do this. I think that starts with the AI PLAN Act.

With that, I yield my time.

Mr. DOWNING. The gentleman yields.

The gentlewoman from Colorado, Ms. Pettersen, is now recognized for 5 minutes.

Ms. PETERSEN. Thank you, Mr. Chairman.

Thank you all for being here today and I am grateful that you all are hosting this opportunity to talk about such an important issue.

I want to address—Mr. Reisman, you talk about how we need to have—it is helpful if we have a regulatory framework at the Federal level, and we could have some complementary State laws as well.

I am really worried because we worked in a bipartisan way to provide a framework to try to continue to move forward here in Congress. We are not known for being the most efficient body in government and our inability to come together to produce a regulatory framework on AI is concerning, when I look at the risks that are involved.

There is great opportunity for efficiency and a lot of promise in all of the different sectors, but I am especially worried in the financial sector and the vulnerabilities that we have.

We know that the risks that already—that we are facing are going to get exponentially worse with AI. Already, 50 percent of fraud has been identified as using AI. Ninety-two percent of the financial institutions surveyed indicate that fraudsters are using generative AI and 44 percent of financial professors report that deepfakes are used in fraudulent schemes.

So, when I think about the impact to our financial system and the gaps that we are leaving here, what are the most important steps? This is opened up for all of you. How do we provide to ensure that we do not have the gaps in supporting our smaller financial systems, so our smaller banks, making sure that they have access to the technology and what would you recommend on providing that framework and those guardrails at the Federal level?

That is a lot to ask, so do not feel the pressure but if anyone wants to add on just some of the key things that we need to be thinking about right now.

Ms. TURNER LEE. I will start.

Ms. PETTERSEN. Okay.

Ms. TURNER LEE. I definitely think the conversation has centered around really advancing national data privacy standards and I think that conversation, as someone has already mentioned, will sort of ring in the beast of data that is actually fueling these systems. If we can come to some bipartisan support on that, I think that would be helpful.

I also think that approaching, Congresswoman, as you said, some of these very nuanced areas, like deepfakes—data is out there that seniors are thinking that their grandchildren are calling them for money. When a senior's economic viability is compromised, that is worse for the whole family. I think going into those verticals where there is bipartisan agreement is another area—we have seen that with the Tools to Address Known Exploitation by Immobilizing Technological Deepfakes on Websites and Networks (TAKE IT DOWN) Act—in terms of anything that has to do with any manipulated content, I think is a step forward.

But I do also think that we have to reimagine what it looks like in terms of regulatory framework that has some guidance and guardrails, which I would be happy to share with your office in more detail.

Ms. PETTERSEN. That is great. I look forward to working with you on this and we have a bill, as well, that I have worked with Representative Flood on; it is the Preventing Deep Fake Scams Act.

This is important because our agencies are only able to take on what we are directing them to do and so, when we look at trying to change with the times to make sure that they have the support they need to address the challenges that we are facing now, can you speak to the greater coordination needed between the regulators, the industry, and the subject-matter experts to develop strategies to protect financial institutions and consumers from fraud using AI?

Ms. TURNER LEE. I sure can, if no one else will.

I think on the deepfake side, I think the challenge that we have is that, because we have multiple agencies with jurisdiction over this particular area, we have not come up with some shared values and goals on, one, how are we defining this.

We have had a lot of conversation, to your point, where we have not necessarily agreed on, on the copyright area, this digital provenance. The same thing, I think, is actually going to go into the deepfake space, where we are looking at extracted video, text, and audio and trying to determine, where is this coming from? That is very hard for us to track at this moment but also coming up with some shared values and goals across the various agencies that are responsible for enforcement over that.

I think if we are able to actually get away from some of that fragmentation, we could actually push forward with really good legislation that protects people from clones.

Ms. PETTERSEN. Great.

Mr. Reisman?

Mr. REISMAN. Yes. I just wanted to say—to praise, in the bill you have called for specifically bringing together industry and expert stakeholders to talk with regulators and government to try and fig-

ure out how to solve these problems together. It has never been more important, as the technology moves quickly, to have all of these voices in the room to problem-solve together.

Ms. PETTERSEN. It seems like this should just—we should bring bills to do this for every agency across—sorry. Thank you. I yield my time.

Mr. DOWNING. The gentlewoman yields.

The gentleman from South Carolina, Mr. Timmons, is now recognized for 5 minutes.

Mr. TIMMONS. Thank you, Mr. Chairman.

Thank you to the witnesses for being here today.

I am pleased that the subcommittee is turning its attention to the use of artificial intelligence in financial markets. Like stablecoins and market structure, this is a technology that will have a significant impact on how American companies and consumers engage with our evolving financial systems.

In order for the United States to remain a global leader in innovation, we must establish clear rules of the road that allow the free market to thrive, and we must do so thoughtfully and effectively.

As I continue to meet with industry leaders, I am increasingly concerned about the growing number of conflicting State laws related to AI. These laws often vary in scope, definition, and enforcement, creating a complex regulatory environment for businesses to operate across multiple jurisdictions.

This patchwork of regulation makes it more difficult for companies to scale, innovate, and remain compliant, while also increasing the risk of inconsistent protections and outcomes for consumers. Without a unified approach, we risk slowing progress and creating barriers that disadvantage both American businesses and the people they serve.

Mr. Gorfine, States have introduced more than a thousand laws related to the use of artificial intelligence. From your perspective, what are the most pressing risks of a State-by-State regulatory patchwork for both consumers and innovators?

Mr. GORFINE. It is a good question.

I want to start by being very clear that when I look at this issue I am talking about it within the context of the regulated financial services sector. There may be issues outside of financial services where States are engaging on AI, but I think that what is really important is to recognize that there are existing Federal and State laws, there is existing regulation, and there is existing guidance that financial institutions adhere to where a patchwork of State laws can absolutely interfere, conflict, or create ambiguity for financial services firms operating on a national level.

I think that is especially clear in the data context. We have existing Gramm-Leach-Bliley Act (GLBA) data privacy laws in place for financial services firms. If you start introducing a patchwork of State approaches, that can be really problematic.

The same goes for when it comes to risk management. There is a very careful risk management framework in place for financial services, and, again, I do worry about interference of State laws there.

So that is how I am thinking about the interplay of State and Federal, especially in the context of the financial services space.

Mr. TIMMONS. I am going to turn that a little bit around and say, while this is going to be difficult for financial services companies to comply and to work across State lines, let us talk about the AI companies that are trying to create these products.

I mean, the Chinese are ahead of us or close to being ahead of us, and the businesses that they have that are focused on this do not have this burden. Would you say that is an equally difficult challenge?

Mr. GORFINE. Yes. I mean, the broader question for AI is that I do think having a proper Federal framework that allows us to operate on a national scale makes really good policy, market, and competitiveness sense. That is something that I would absolutely encourage because, as we are describing, it is not even just a national competition; it is global, right? These activities transcend borders and having a coherent, kind of, Federal framework there makes good sense in the AI context.

Mr. TIMMONS. Thank you for that.

While there are many challenges involved in legislating and establishing clear rules of the road for artificial intelligence, there are also powerful opportunities. Artificial intelligence has the potential to equip financial institutions with advanced tools to better serve their clients, improve efficiency, and reduce risk.

As I mentioned earlier, one of the most common concerns that I hear from both large financial institutions and smaller community banks and credit unions is the burden that regulatory compliance places on their operations. These requirements strain both their financial resources and staffing capacity.

Artificial intelligence offers a promising solution to help automate and streamline compliance processes, allowing firms to redirect time and resources toward innovation and customer service.

Dr. Lau, given the significant burden that regulatory compliance places on financial institutions of all sizes, how is artificial intelligence helping firms streamline these processes and manage risk more effectively?

Mr. LAU. Thank you for the question.

Definitely, we see applying AI to compliance workflows as one of the highest return on investment (ROI) activities that we are seeing banks implement today, including smaller community banks and regional banks, right? The reason is not just so they are able to automate more workflows where you have low staffing, but also they open up new opportunities for greater compliance. Think about continuous monitoring, 24-hour audits, et cetera, right?

On the flip side, you also need to be able to look at, as you are introducing AI to these regulated workflows, what type of guardrails or controls are put around those AIs so that they are able to follow the right policies and procedures when executing those workflows. That, actually, in many cases, is the more challenging problem, is, how do you rein in this autonomous AI to actually carry out the tasks effectively and in compliance?

Mr. TIMMONS. Thank you for that.

Emerging technology has the ability to make sure that the U.S. is the center of the global economy for decades to come, and we have to get this right.

With that, Mr. Chairman, I yield back.

Mr. DOWNING. The gentleman yields.

The gentleman from New York, Mr. Torres, is now recognized for 5 minutes.

Mr. TORRES. Thank you, Mr. Chair.

AI is the most transformative technology of our time. The rise of generative AI could prove to be as revolutionary as the advent of writing or the advent of the printing press. Whether AI will create a better world or a worse world, no one knows for sure. What we do know for sure is that the world will be radically different from anything we have seen before.

There is nothing new about the use of AI in finance. What is new is the use of generative AI, large language models, in finance.

So my first question: What are the capabilities in finance that a large language model can perform that legacy AI has been historically unable to perform?

Anyone who can answer that question is free to do so.

Mr. LAU. I can go ahead and jump in.

I think Dr. Cox actually mentioned at the beginning, the reason why there is so much investment in this space is because it is general-purpose AI, meaning that it could actually do many different things. You can prompt it in infinitely different ways to carry out certain types of workflows that are very specific to your day-to-day, all the way to doing things like continuous monitoring, et cetera, right?

So that is the power of the technology, are these general-purpose AI models—

Mr. TORRES. What is the best new use—what is the best new use case in finance?

Mr. LAU. Yes. I would say—I mentioned compliance as one of the very high ROI activities because of the manual effort involved in that but I would say, the most mature one that is delivering the highest ROI that we see is around developer productivity, the ability for these AI agents to actually go and build applications from scratch, empowering even, we are seeing, small community banks to leverage these to actually really accelerate their own software development and integrating their infrastructure stack.

Mr. TORRES. Can an AI banker outperform a human banker?

Mr. LAU. I think one of the biggest challenges to actually having the AI outperform a human banker is to make sure that it complies with the common sense of a human banker, right? That has actually been challenging when you look at AI agents, to give them that type of common sense out of the box, even given the vast amounts of data that they have been trained on.

So building guardrails that a human is trained to follow is something that is still an open challenge but something we are helping a lot of these banks with today.

Mr. TORRES. Because I know Members of Congress are irreplaceable, but I am wondering if bankers and traders are replaceable by AI.

What is the impact of AI on concrete applications like credit scoring, loan approvals, fraud prevention and detection? Like, does AI lead to more loan approvals or fewer? Does it lead to more inclusive credit scoring or less inclusive credit scoring? What is the outcome so far.

Ms. TURNER LEE. I can—oh, do you want—well, I can point to that.

I mean, I think we thought, because of the objective nature of AI, that it would be easier to see more loan approvals, more credit applications actually confirmed—

Mr. TORRES. Where did we get this notion that AI is objective?

Ms. TURNER LEE. Well, we thought—

Mr. TORRES. The data comes from the internet.

Ms. TURNER LEE. Well, that is what I was going to say. We thought—

Mr. TORRES. It is a reflection of human nature.

Ms. TURNER LEE. That is right and because the data that is fueling the AI systems basically comes from us—and particularly in generative AI, it is curated data. It is not necessarily predictive data. It is data that exists on people on the public internet—we are actually in some cases the same results and in other instances worse.

We have seen that also on the housing appraisal side when it comes to homeownership, that even if you scrub your entire home of any type of remnant or artifact of you, the AI will still generate the same results, because AI uses other proxies—your address, your net worth in your community, et cetera—Congressman.

So I think, when we say that AI is better than humans or a banker can outperform a human, probably in time, but expeditious calculation of how we make these decisions comes at foreclosing on opportunities economically for various consumers.

Mr. TORRES. Can AI—I guess I was going to ask, can AI be harnessed to expand access to capital, to expand access to credit? One of my frustrations—I have real frustrations with traditional credit scoring methodology. Can AI detect new patterns of evaluating creditworthiness?

Mr. GORFINE. I think that is right. I mean, traditional credit scores are highly correlated with protected class characteristics.

One thing I would suggest is to consider second-look applications of AI. What that means is, if you take an existing decline pool and you run the decline pool through, kind of, cutting-edge, gen-AI-related underwriting models, you will at worst case result in another decline, but you may actually start pulling some approvals through that process, and you mitigate some of your initial concerns around the impact of such models.

So I think there are ways to smartly start testing these models in a way that upholds fairness.

Mr. TORRES. The question for me—because, inevitably, AI is going to have some measure of algorithmic bias—

Ms. TURNER LEE. That is right.

Mr. TORRES [continuing]. right? The question for me is not whether it is completely free of bias but whether we can make it less biased than the human alternative.

Is that an achievable mission?

Ms. TURNER LEE. Well, I think it is achievable if there is transparency, first and foremost, that an AI model is making the decision on behalf of the financial—on a financial application. Many people do not know that AI is actually being used to make those credit decisions, so you have to start with public disclosure.

I think the second thing, to your point, is, we do need a stat that is able to actually evaluate what the disproportionate impact—disparate impact is.

Mr. TORRES. I am about to be replaced by AI but thank you.

Mr. DOWNING. The gentleman's time has expired.

I would like to thank all the witnesses for your testimony today.

Without objection, all members will have 5 legislative days to submit additional written questions for the witnesses to the chair. The questions will be forwarded to the witnesses for their response.

Witnesses, please respond no later than October 23, 2025.

[The information referred to can be found in the appendix.]

Mr. DOWNING. With that, this hearing is adjourned.

[Whereupon, at 3:47 p.m., the subcommittee was adjourned.]

APPENDIX

MATERIALS SUBMITTED FOR THE RECORD



Stay updated with the latest news. Allow notifications from Daily Caller.

Allow



REP WARREN DAVIDSON: Privacy, Artificial Intelligence, And Congress

DAILY CALLER

Justin Sullivan/Getty Images

September 17, 2025 1:45 AM ET

Rep. Warren Davidson Contributor

As Artificial Intelligence (AI) continues to grow in both popularity and overall ability, Congress needs to act proactively — not reactively — to protect the American people by passing comprehensive frameworks to address concerns with AI. Recently **named** one of Time Magazine's 100 Most Influential Voices on AI, Pope Leo XIV recently expressed similar concerns and spoke of the need to "use these gifts [of technology] wisely, ensuring they serve the common good." What's more, a new **poll** from Pew Research shows that over 53% of Americans say AI is doing more to hurt than help people keep their personal information private.

Unfortunately, Congress hasn't taken AI regulatory issues seriously — including the privacy implications. Many of the worst risks with AI don't come from the technology alone but from the absence of strong privacy safeguards. Without clear limits on how data is collected, stored, and shared, AI becomes a powerful tool for exploitation and surveillance. That's why any serious conversation about Artificial Intelligence begins with privacy.

Privacy forms the base layer for ethical Artificial Intelligence. How do you safeguard your own data? Who is liable if your data is compromised, especially by AI? Does simply entering a query into a search engine or engaging in dialogue with AI void a reasonable expectation of privacy? Can the company then claim that data as its own? Can that search or query be made public? Sold and



Stay updated with the latest news. Allow notifications from Daily Caller.

Fast. School Districts Aren't Ready).

These are important questions that Congress must answer, and if Congress is unable to answer them, then state governments may have to take the lead for now. In the meantime, surveillance capitalism and government spying on its own citizens has run amok. At the intersection, Palantir has been commissioned to develop an AI tool to unify the data in every federal database, turning it into useful, easily accessible information.

The Patriot Act massively expanded domestic surveillance. The Bank Secrecy Act ended any claim to privacy in your financial dealings. Most Bureaus of Motor Vehicles monetize the data that citizens are required to provide to drive or obtain a Real ID. Now the government is buying data that would otherwise require a warrant or subpoena — circumventing the Fourth Amendment.

While the original House version of the Big Beautiful Bill included a 10-year moratorium on state or local regulation of AI, the Senate wisely voted 99-1 to remove this provision. The Senate's removal of the AI provision influenced my decision to support the final passage of the bill. Thankfully, stripping the AI regulatory exemption from the Big Beautiful Bill offers hope of accountability.

Freedom surrendered is rarely reclaimed. Congress may decide to keep things the way they are now with respect to privacy, but they shouldn't. The Fourth Amendment doesn't say, "If you have nothing to hide, you have nothing to fear." Instead, it says the government requires probable cause and a warrant or subpoena to obtain even limited ability to search or seize your property or information. We desperately need to restore a government small enough to fit within the Constitution.

In 2003, the movie *iRobot* laid out three laws designed in a fictitious future (2035) to safeguard human interaction with their Artificial Intelligence (AI)-controlled



Stay updated with the latest news. Allow notifications from Daily Caller.

1. A robot may not injure a human being or, through inaction, allow a human being to come to harm.
2. A robot must obey the orders given it by human beings except where such orders conflict with the First Law.
3. A robot must protect its own existence as long as such protection does not conflict with the First or Second Law.

With respect to commerce, privacy resolves the fundamental issue of who owns what data. With that resolved, Artificial Intelligence remains challenging. *iRobot* wasn't far off the mark — the First Law says no human should come to harm. Today, we lack even that baseline for AI. Perhaps we should add a few rules to that list, but we must define the parameters because the real-world consequences are already upon us.

Hollywood imagined rules in 2003 to keep humans safe from AI. Yet, over two decades later, Congress still has not written any rules to do the same. Even worse, many of my colleagues seem entirely unconcerned with the potential repercussions that could soon become reality. Failure to act is its own decision, but AI is moving far faster than Congress and momentum in the wrong direction makes change harder as time passes. If we don't act swiftly, our current understanding of what "privacy" means could become a relic.

Warren Davidson represents Ohio.

The views and opinions expressed in this commentary are those of the author and do not reflect the official position of the Daily Caller News Foundation.

All content created by the Daily Caller News Foundation, an independent and nonpartisan newswire service, is available without charge to any legitimate news publisher that can provide a large audience. All republished articles must include our logo, our reporter's byline and their DCNF affiliation. For any questions about our guidelines or partnering with us, please contact licensing@dailycallernewsfoundation.org.



**America's
Credit Unions**

Jim Nussle
President & CEO
202-508-6745
jnussle@americascreditunions.org

99 M Street SE
Suite 300
Washington, DC 20003

September 17, 2025

The Honorable Bryan Steil
Chairman
Subcommittee on Digital Assets,
Financial Technology, and
Artificial Intelligence
Committee on Financial Services
U.S. House of Representatives
Washington, DC 20515

The Honorable Stephen Lynch
Ranking Member
Subcommittee on Digital Assets,
Financial Technology, and
Artificial Intelligence
Committee on Financial Services
U.S. House of Representatives
Washington, DC 20515

Re: Tomorrow's Hearing: "Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness"

Dear Chairman Steil and Ranking Member Lynch:

On behalf of America's Credit Unions, I am writing regarding the Subcommittee's hearing entitled "Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness." America's Credit Unions is the voice of consumers' best option for financial services: credit unions. We advocate for policies that allow the industry to effectively meet the needs of their over 144 million members nationwide.

America's Credit Unions appreciates the Subcommittee holding this hearing on the potential transformative properties of artificial intelligence (AI) and the role of regulators in overseeing its deployment. Credit unions are engaging in collaborative innovation by working with credit union service organizations (CUSOs) and other financial technology vendors to leverage the power of AI safely and securely to help meet the financial needs of their members. From managing internal procedures to AI chatbots, to underwriting assistance and fraud detection—credit unions are seeing firsthand how AI is increasing staff efficiency, automating previously laborious tasks, reducing paperwork, and expediting loan decision making processes.

AI-powered fraud analytics can enhance credit union efforts to prevent financial crime by improving detection of irregular financial behaviors. Many credit unions are already using third-party technology bundled with debit and credit card products to prevent fraudulent transactions or to flag suspicious transactions. Some credit unions are also using AI to assist with check fraud, particularly by identifying check washing schemes. In some cases, this technology leverages AI and machine learning processes (e.g., neural networks) to develop predictive models for fraud mitigation purposes. Credit unions are eager to adopt more effective fraud management tools given the increasing prevalence of card-not-present fraud and the impossibility of manually monitoring transaction patterns.

September 17, 2025
Page 2 of 4

AI is already providing value in improving Know Your Customer identity verification processes and has the potential to reduce Bank Secrecy Act (BSA) and anti-money laundering (AML) compliance costs by minimizing the burden of filing suspicious activity reports. Similarly, AI could be used to satisfy regulatory notification standards if an institution experiences a reportable cyber incident. However, for these benefits to be fully realized, regulators would need to signal greater acceptance of AI for regulatory reporting purposes.

As the financial services industry continues to rapidly innovate and further integrate AI tools into their core systems, credit unions have a variety of concerns about how their regulators approach AI and the potential risks AI poses to systems they use to serve their members.

America's Credit Unions Supports Efforts to Reduce Barriers to Innovation

A stringent stance from regulators on AI usage in the financial sector could disproportionately harm credit unions and smaller institutions while benefiting the largest incumbents. It is reasonable to expect explanations about the outputs and decisions an AI system can generate. However, demanding a detailed inspection of source code, algorithmic weights, or other technical aspects underlying AI models would impose significant and impractical burdens on credit unions that rely on proprietary models from third-party vendors. Such requirements could divert scarce in-house developer resources towards compliance activities, which larger institutions can more readily afford.

Past statements from the Consumer Financial Protection Bureau (CFPB) suggesting that AI algorithms function as opaque “black boxes” introduce uncertainty about the level of regulatory scrutiny required for algorithmic processes, especially in the absence of evidence of harm to consumers. This skepticism puts credit unions and other community institutions at a significant disadvantage. Instead of being able to innovate and explore new technologies like AI freely, credit unions may find themselves compelled to demonstrate that these innovations pose no risk—a daunting task. When regulatory discussions approach technology in this manner, it reinforces the idea that smaller entities could face obstacles to success, as only the largest institutions may possess the resources necessary to meet such rigorous risk management standards.

Regulatory assessments of new technologies must embrace balanced and flexible approaches to risk management that can accommodate innovation while protecting consumers. Often this balance is expressed as “responsible innovation,” which regulators have sometimes supported through official statements and relevant guidance. For example, in 2019, an interagency statement issued by the National Credit Union Administration (NCUA) and other banking regulators recognized that the “use of alternative data may improve the speed and accuracy of credit decisions and may help firms evaluate the creditworthiness of consumers who currently may not obtain credit in the mainstream credit system.” This statement has helped legitimize the use of new data for credit underwriting purposes while simultaneously acknowledging that “the use of certain alternative data may present no greater risks than data traditionally used in the credit evaluation process.” While alternative data is an input rather than a new technology

September 17, 2025
Page 3 of 4

unto itself, the interagency statement demonstrates how consensus within the regulatory community can help reduce uncertainty and drive innovation forward in a way that is beneficial for consumers.

Additionally, financial sector regulators should tailor future actions related to AI in a way that distinguishes between the use of AI technology by regulated versus unregulated institutions. For highly regulated financial institutions, such as credit unions, an appropriate regulatory framework should recognize the need for less prescriptive intervention and greater accommodation of innovation through pilot programs, no-action letters, waivers, and elimination of outdated rules. In an environment where non-bank fintech companies may enjoy a less rigorous supervisory oversight than traditional financial institutions, regulators should be exploring frameworks that make innovation accessible not just to the largest and most sophisticated entities, but also to smaller, community-based institutions. The need to establish a fair playing field cannot be overstated. Additionally, an overly complex or intrusive supervisory framework for assessing the inner workings of AI algorithms would likely deter credit unions from investing in these technologies and frustrate efforts to partner with CUSOs and other technology providers.

Third-party providers of AI services may not grant customers access to the proprietary code that underlies machine learning algorithms or models. In these circumstances, demanding that end-users at financial institutions attest to the specific operational parameters of an AI product would not be feasible and would likely inhibit the use of such technology by credit unions unable to develop these products in-house. While AI explainability may be a key consideration for regulators, innovation cannot reasonably take place in an environment where every line of code is scrutinized and financial institutions must prove a negative: that a model is incapable of error. A more reasonable approach would be to encourage regulators to understand how AI risks can be addressed through the application of existing laws and regulations, which are not limited to any particular mode of decision making or technology. The evaluation of AI outputs using existing compliance procedures can identify and prevent consumer harm before it occurs.

Ensuring Fairness in Lending with AI Optimization

Credit unions were created to offer provident credit to all members of their communities, and this organizing principle helps to explain the prevalence of robust relationship lending models across the industry. As cooperatives that are directly accountable to their member-owners, credit unions are focused on developing long-lasting, trusted relationships—an interest that is best served by adhering to core principles of equality and fairness.

Credit unions follow existing regulations and guidance implementing the Equal Credit Opportunity Act (ECOA) and other anti-discrimination laws and have a track record of exceptional fair lending compliance. Most credit unions engage in self-tests or self-evaluations as part of their ongoing monitoring of fair lending risks. While self-evaluations can vary in terms of their scope and sophistication based on a credit union's risk profile, they generally encompass review of denied applications, comparisons of loan files, analysis of Home Mortgage Disclosure

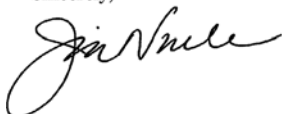
September 17, 2025
Page 4 of 4

Act (HMDA) data, and review of lending policy exceptions. Self-tests can be similarly varied and encompass a variety of analytical techniques (e.g., surveys, use of test applicants, review of credit transaction records). Both types of testing could function as post-hoc methods for evaluating the results of AI-driven lending decisions. These existing methods should be regarded as effective for detecting and addressing fair lending risks. Deconstructing the entirety of an AI algorithm to address explainability or overfitting risks would be costly and less productive for examination purposes.

AI and machine learning can improve access to credit and access to housing. Using automated underwriting as part of partnerships with AI fintech companies has helped a number of credit unions increase approval rates and get to “yes” even more often for minority borrowers while holding the line on delinquencies. These credit unions prioritize fair lending by ensuring a human verifies any denials determined using automated underwriting tools. In short, AI can help reduce the racial homeownership gap when the technology works as intended with proper human oversight. Because credit unions support legislative and regulatory initiatives that promote the availability of credit to all creditworthy applicants, they recognize that responsible use of AI mandates compliance with existing applicable law. That said, credit unions and America’s Credit Unions expect that legislators and regulators will weigh the regulatory burden on credit unions against the benefit to consumers when implementing new laws and regulations that govern the use of AI in lending and housing. To do otherwise may chill technological progress and the benefits that may accrue to consumers through the responsible use of AI.

On behalf of America’s credit unions and their over 144 million credit union members, thank you for the opportunity to share our views. We look forward to continuing to work with you on these important issues. Should you have any questions or require any additional information, please contact me or Greg Mesack, America’s Credit Unions’ Senior Vice President of Advocacy, at gmesack@americascreditunions.org.

Sincerely,



Jim Nussle

cc: Members of the Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence



September 30, 2025

The Honorable French Hill
Chair
Committee on Financial Services
U.S. House of Representatives
Washington, D.C. 20515

The Honorable Maxine Waters
Ranking Member
Committee on Financial Services
U.S. House of Representatives
Washington, D.C. 20515

Re: NFHA Comments on H.R. 4801, *Unleashing AI Innovation in Financial Services Act*

Dear Chair Hill and Ranking Member Waters:

The National Fair Housing Alliance® (NFHA™) submits these comments concerning H.R. 4801, the *Unleashing AI Innovation in Financial Services Act*.

NFHA is dedicated to eliminating discrimination in housing and ensuring equal opportunity in all housing-related financial services. We acknowledge the importance of innovation, including the use of artificial intelligence (AI), in financial services. However, we are deeply concerned about proposals that create regulatory safe harbors or sandboxes that could permit or facilitate civil rights or consumer protection violations, particularly through automated systems and “black box” AI models in financial, housing, and related services.

Summary of NFHA Position

- NFHA is generally opposed to sandboxes that excuse or weaken enforcement of civil rights and consumer protection laws, especially in high-risk¹ financial and housing contexts.
- NFHA vehemently opposes any legislation that excuses, suspends, or gives safe harbor from liability for violations of civil rights or consumer protections.
- NFHA opposes the use of No Action Letters for high-risk AI models, such as those determining credit, pricing, loan approval, or access to housing.
- NFHA supports the responsible use of sandboxes or innovation labs, with appropriate protections and remedial provisions, and supports the utilization of these tools to test solutions that benefit humanity and spur economic mobility and growth.

¹ Article 6(2) of the EU AI Act defines high risk AI systems as those that pose a significant risk to human health, safety or fundamental rights, and are subject to extensive obligations and requirements such as critical infrastructure, education admissions, employment recruitment, creditworthiness and public services, law enforcement and justice, migration and border control and the administration of justice.



Key Concerns with H.R. 4801

As introduced, H.R. 4801 would:

1. Allow financial institutions to experiment on people without expectation of enforcement actions

NFHA is concerned that creating AI Innovation Labs that insulate entities from enforcement actions undermines the most fundamental safeguard of civil rights and consumer protection laws. AI has the capacity to create, perpetuate, amplify, and accelerate historical patterns of discrimination in housing and financial services. Allowing firms to experiment on people without the expectation of enforcement invites high-risk trials in financial and housing markets where discrimination has already produced unfair outcomes and severe racial homeownership and wealth gaps. Moreover, NFHA stresses that regulators must define *model risk* to include discriminatory outcomes for consumers, not just financial loss to firms. By suspending enforcement during the experimentation, H.R. 4801 risks creating “testing grounds” where discriminatory models can subsequently harm people on a prohibited basis (e.g., race, religion and gender) with little or no recourse.

2. Permit industry to design their own alternative compliance strategies for civil rights and consumer protection regulations

NFHA strongly opposes any provision allowing regulated entities to “comply differently” with civil rights and consumer protection laws and design their own “alternative compliance strategy.” History demonstrates that alternative compliance schemes in housing and credit markets often mask or reproduce discrimination—whether through redlining, racial steering, fomenting segregation, disparate loan pricing, or biased scoring models. Our nation’s civil rights laws and consumer protections exist because of our long history of race-based, race-conscious laws and policies, many perpetuated by the federal government, that have unfairly denied people the right to housing, financial services, and other opportunities. These race-conscious laws, policies, and practices created unfair constructs that still exist today: restrictive zoning, residential segregation, biased models, etc. These unfair systems still cause discrimination in our markets. This fact cannot be ignored or overlooked in a regulatory sandbox environment. Rather, a regulatory sandbox must account for and protect against both intentional and unintentional discrimination.

Legal tools and protections, such as disparate impact, exist precisely because facially-neutral policies, models, and systems can yield unlawful outcomes that can disproportionately harm people on a prohibited basis. Thus, a thoughtless housing practice can be as unfair to underserved and/or marginalized people as a willful scheme. Permitting regulated entities to experiment on people with modified obligations without rigorous civil rights guardrails and consumer safety standards could erode the enforceability of the Fair Housing Act, Equal Credit Opportunity Act, and related statutes, and undermine Congress’ carefully constructed statutory schemes. In effect,



each financial institution would be able to write its own laws under the bill's current requirements for alternative compliance plans.

3. Provide insufficient protections for financial institutions that want to experiment on real people in real-world scenarios

The protections in the EU AI Act can be used as a starting point. Unlike the EU AI Act, this bill does not contain specific protections for an AI Test Project that is conducted in real-world scenarios with real people. For example, the EU AI Act requires:

- Time limits. The testing lasts no longer than necessary to achieve its objectives and no longer than six months.
- Deployer responsibilities. When AI System Providers and Deployers collaborate on testing, Deployers must be informed of all aspects of the testing that are relevant to their decision to participate and given relevant instructions for use. The AI System Provider and Deployer must agree on their roles and responsibilities to meet testing requirements.
- Consumer consent. The subjects (e.g., consumers) of the testing must provide informed consent, which requires, among other things, that they have been given information about their rights and the nature and objectives of the testing; any possible inconvenience the testing may cause; and the conditions under which the testing will be conducted.
 - Opt-out. Any subjects of the testing may, without any resulting detriment and without having to provide any justification, withdraw from the testing at any time by revoking their informed consent and may request the immediate and permanent deletion of their personal data.
 - Protection of vulnerable subjects. Subjects of the testing who are vulnerable persons due to their age or physical or mental disability are appropriately protected.
- Qualified oversight. The testing is overseen by the AI System Providers and Deployers through persons who are suitably qualified and have the necessary capacity, training, and authority to perform their tasks.
- Agency authority. Agencies can require AI System Providers to supply information, carry out unannounced remote or on-site inspections, and perform checks on the development of the testing and the related products to ensure the safe development of testing.
- Incident reporting. The AI System Provider must report to the Agency any serious incident identified during the testing and adopt immediate mitigation measures or, failing that, suspend the testing until such mitigation takes place, or otherwise terminate it.
- Continued liability. The AI System Provider remains liable to consumers and others for any damage caused during the testing.

Moreover, while NFHA recognizes the value of requiring detail in the applications, the proposed criteria are also insufficient because they do not mandate civil rights or impact assessments. Our nation's civil rights laws are central to ensuring fairness as well as safety and soundness in housing and financial services markets. These protections cannot be protected without data on race, ethnicity, and other protected characteristics. An AI Test



Project application that fails to examine disparate impacts on communities of color, and other protected groups is incomplete, misleading, and can generate harmful—and sometimes irreversible—outcomes. Similarly, termination dates and economic impact analyses cannot substitute for mandatory evaluations of bias, discrimination risk, and less discriminatory alternatives. Without explicit requirements to measure and disclose potential unfair or biased outcomes, these application provisions risk producing paperwork exercises that legitimize harmful experiments on people rather than preventing them.

4. Provide insufficient real-time monitoring requirements and agency oversight

NFHA is concerned that while annual reporting requirements create a veneer of transparency, they are retrospective and may not prevent harm to people in real time. NFHA has cautioned Congress and regulators that relying on *ex post* reporting has historically failed to stop crises such as the 2008 predatory lending and foreclosure crisis that devastated neighborhoods and our economy. Annual reports to Congress will not provide meaningful protection to borrowers who may already have been denied loans or housing, steered into higher-cost credit or housing, or targeted for discriminatory terms and pricing during a “test project.” Moreover, the bill does not require the reports to be released to the public and requires that the names of the financial institutions be suppressed.

In addition, unlike the EU AI Act, this bill does not require robust agency oversight. The EU AI Act is markedly different because it creates a framework for give and take between the agency and the regulated entity. The EU AI Act requires the agency to provide guidance to the regulated entity throughout the process regarding how to comply with the law. By contrast, HR 4801 simply sets the regulated entity free after agency approval of the application. It is imperative that the federal financial regulators engage in robust supervision and enforcement activities during the life of AI models, not simply observe and document outcomes in vague annual reports. Moreover, public reporting must mandate data on impact assessments, including disaggregated demographic impacts; otherwise, Congress and the public will be unable to evaluate whether the AI Test Projects advance housing opportunity or perpetuate discrimination.

5. Contain many technical issues, creating an incredibly complex framework with little real regulatory relief for financial institutions

HR 4801 contains many technical issues, which ultimately mean that the bill neither provides adequate protections for consumers nor provides real regulatory relief for financial institutions. The biggest issue is that the bill does not account for the complex regulatory and enforcement structure for financial institutions in the United States. For example:

- The statute only exempts the regulated entity from enforcement actions by a financial regulatory agency, not from an action by another federal enforcement agency (e.g., DOJ, FTC), a state agency, or a private plaintiff.



- A financial regulatory agency may still take enforcement action against a regulated entity with respect to “fraud” or for engaging in an “unsafe or unsound practice” related to the AI Test Project.
- A financial regulatory agency may file a civil action in an appropriate U.S. district court to enjoin an AI Test Project if the agency determines that the AI Test Project—
 - Presents an “immediate danger” to consumers or investors; or
 - Presents a “risk” to financial markets, to the federal deposit insurance fund, of a violation of AML laws, or to national security.
- Note: The terms “fraud,” “unsafe or unsound practice,” “immediate danger” and “risk” are not defined and could vary greatly from administration to administration as well as regulator to regulator.

Finally, the bill suffers from a technical issue in who is provided with the limited legal exemption. The bill applies to “regulated entities,” which are financial institutions, but does not explicitly cover AI providers. Similarly, an “AI Test Project” is defined as “a financial product or service that (A) falls under the jurisdiction of a financial regulatory agency, (B) makes substantial use of AI, and C) is, or may be, subject to a Federal regulation or Federal statute.” Generally, it is the financial institution that falls under the jurisdiction of the agency, rather than the product or service. These technical issues would need to be addressed if the bill hopes to provide even a modicum of regulatory relief.

NFHA Recommendations for Amendments

To make H.R. 4801 compatible with civil rights and consumer protections such as anti-discrimination laws and prohibitions on Unfair, Deceptive or Abusive Acts or Practices, NFHA recommends the following amendments, at a minimum:

1. **Explicit civil rights / fair housing requirements**
 - Require all AI Test Project applications to include an *impact assessment* centered on *civil rights and fair housing principles*.
 - Mandate *independent, external bias audits, including disparate impact findings* before, during, and after projects to build best practices.
2. **No safe harbor from liability for antidiscrimination laws**
 - Make explicit in **Section IV (Rules of Construction)** that nothing in this Act alters obligations under the Fair Housing Act, Equal Credit Opportunity Act, Civil Rights Act, or other federal anti-discrimination laws.
3. **Restrict high-risk AI Test Projects proposed to be performed on real people in real-world scenarios**
 - Prohibit or strictly condition AI Test Projects in high-risk areas (e.g., credit underwriting, loan pricing, mortgage approval, housing access) where bias and discrimination risks are greatest.



- Requirements should include:
 - Consumer consent and right to opt out.
 - Strict time limits, such as a minimum evaluation period of one year and a maximum of two years — particularly AI systems involving predictive modeling or consumer risk assessment, or relative to the time window used to set the AI model objectives.
 - Incident reporting to the agency.
- 4. Provide conditional relief for AI Test Projects experimenting with fairness safeguards**
 - Explicitly authorize AI Test Projects for testing techniques that diminish discrimination in financial services, including housing and consumer lending; mitigate damage caused by climate change or environmental injustice; and mechanisms that provide cures for disease or advance public health.
 - Require strict guidelines: transparency, consumer disclosures, accountability, independent audits, participation of civil rights and consumer advocates, and clear redress mechanisms for any harmed individuals.
- 5. Further processing of personal data and demographic information:** As an exception to privacy laws, personal data lawfully collected for other purposes may be processed solely for the purposes of developing, training and testing certain AI systems under the AI Test Project under certain conditions and with meaningful protections for people, including that the AI system is being developed to advance fair housing and lending opportunities.
- 6. Establish realistic timeframes for AI Test Project evaluation:** It is recommended that the bill include a minimum evaluation period that ensure regulatory timelines are aligned with industry practices and technical realities, particularly for AI systems involving predictive modeling or consumer risk assessment.
 - **Rationale:**
 1. AI models in financial services often rely on historical data windows to define their target variables or labels — technical terms referring to the outcomes the model is trained to predict (e.g., loan default, fraud detection).
 2. These targets are typically derived from one- to two-year data windows, reflecting industry norms for model calibration and performance benchmarking.
 3. Imposing a shorter regulatory review period may lead to premature conclusions, misaligned oversight, and ineffective policy enforcement.
 - **Policy implication:**
 1. A one- to two-year timeline aligns with the lifecycle of most financial AI models and ensures that regulatory assessments are based on statistically meaningful results.
 2. Flexibility may be considered for high-risk systems or those with shorter operational cycles, but the default standard should reflect prevailing industry practices.



- **Note:** While terms like “target” and “label” are technical, they can be referenced in legislative drafting with explanatory footnotes or reserved for internal regulatory guidance documents.
7. **Comparative justification with EU AI Act Innovation provisions:** It is recommended to strengthen the rationale for benchmarking the bill against comparable innovation-related provisions in the EU’s AI Act. This clarification and comparative framing serve two strategic purposes:
- **Global regulatory alignment:** A significant number of U.S.-based financial institutions operate across EU jurisdictions and deploy AI systems that are often built on shared architectures. While data inputs may vary by region, the underlying models and risk profiles are frequently similar. Aligning innovation safeguards and incentives with EU standards can reduce compliance friction and promote cross-border interoperability.
 - **Competitiveness and innovation incentives:** The EU AI Act includes targeted provisions to support innovation, particularly for smaller companies and research institutions. By referencing these measures, the bill can demonstrate a commitment to fostering responsible AI development while maintaining global competitiveness — a priority for both industry stakeholders and national economic strategy.
8. **Independent oversight and public accountability**
- Require robust federal agency oversight and guidance.
 - Provide the funding and resources needed so that Agencies can employ, train, and retain the staff needed to conduct robust analyses of AI Test Projects within the specified timeframe.
 - Require financial institutions to retain civil rights and consumer protection experts in the design and evaluation of AI Test Projects.
 - Require financial institutions to retain the services of independent, third-party AI system auditors.
 - Mandate public disclosure of the specific financial institution’s AI Test Project rationales, safeguards, and outcomes.

Conclusion

H.R. 4801, as drafted, risks undermining civil rights and consumer protections in the financial services marketplace. NFHA cannot support the bill in its current form. However, we would welcome a reworked version that:

- preserves full accountability under civil rights and consumer protection laws;
- restricts sandboxes in high-risk contexts; and
- explicitly supports sandboxes designed to advance access to economic opportunities and benefit the public interest, such as reducing discrimination in financial services and spurring economic mobility—under strict guidelines and safeguards.



NFHA stands ready to work with the Committee to ensure that innovation in AI strengthens rather than undermines the rights and protections that are fundamental to our democracy.

Sincerely,

A handwritten signature in black ink, which appears to read "Lisa Rice". The signature is fluid and cursive, written in a professional style.

Lisa Rice
President & CEO
National Fair Housing Alliance



215 Pennsylvania Avenue, SE • Washington, D.C. 20003 • 202/546-4996 • www.citizen.org

September 22, 2025

The Honorable French Hill, Chair
 The Honorable Maxine Waters, Ranking Member
 Honorable Members
 House Financial Services Committee
 Washington, DC 20515

Re: AI Sandbox Immunity hearing and legislation

Dear Committee members,

On behalf of more than one million members and supporters of Public Citizen, we vigorously oppose Chair Hill's [H.R. 4801, the Unleashing AI Innovation in Financial Services Act](#), better described as "sandbox" get-out-of-jail-free legislation. While AI may help make financial services more efficient, case studies to date reveal serious dangers—from discriminatory lending and predatory upselling to hallucinations and opaque decision-making. This bill would allow Wall Street and Big Tech to rewrite their own compliance rules, granting automatic approval if regulators fail to act on time and shield firms from accountability. It weakens oversight by barring other regulators from stepping in and anonymizing reporting, leaving the public and Congress in the dark about which firms are experimenting on consumers.

Instead of providing immunity from long-standing consumer protections, Congress should require rigorous beta testing of AI systems under pre-screened conditions, with independent audits, frequent review, and full transparency. Financial innovation should never come at the cost of consumer safety, market stability, or democratic accountability. Moreover, allowing the financial industry to receive waivers from federal laws and regulations, while simultaneously allowing them to dictate the rules they feel like following is an insult to the American people and Main Street.

We note that the hearing where Chair Hill introduced this bill demonstrated the dangers of AI; the hearing certainly failed to support the concept of an immunity sandbox. Both Democrat and Republican lawmakers expressed well-founded concerns.

For questions, please contact Bartlett Naylor at bnaylor@citizen.org, and/or JB Branch at jbranch@citizen.org

Sincerely,

Public Citizen



DynamoFL, Inc.
 2261 Market Street #4629
 San Francisco, CA 94114
 E: Christian@dynamo.ai

Dr. Christian Lau Responses to Questions for the Record

Questions for the Record for Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence of the Committee on Financial Services

Hearing: Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness

From Rep. Rose:

1. Dr. Lau, your prepared testimony highlights that AI can provide more scalable customer service. While it is clear that AI has the potential to assist consumers effectively with routine service issues, there are unique, complex, or stressful situations where human intervention is invaluable. Do you believe AI will eventually replace customer service representatives entirely in the financial services industry, or do you think human representatives will continue to play an essential role?

AI is already valuable for handling routine, high-volume service needs, where automation improves speed and consistency. But in financial services, many interactions involve regulated advice, sensitive data, or decisions with consequential impacts. In these situations, human oversight is critical.

As AI systems become more capable and helpful, the role of humans will evolve. The essential function may shift from directly providing every answer to setting direction, supervising outputs, and ensuring the controls around these systems work as intended. That means humans must be able to intervene, test, audit, and assess the quality of AI-generated guidance. With the right guardrails, this can expand access to timely support while keeping consumer protection front and center.

At the same time, it's important that the U.S. continues to strengthen both its highly skilled AI workforce and its broader service workforce. As lower-skilled tasks become more automated, the real value shifts toward AI governance, oversight, and effective use of these tools — areas where human judgment is indispensable. Upskilling workers to use, supervise, and govern AI systems not only protects jobs; it positions American workers and institutions to stay at the forefront of global competitiveness.

So while we may see more AI-driven representatives over time, humans will remain essential — not necessarily for every interaction, but for the oversight, governance, and judgment that ensure these systems operate safely, fairly, and in a way that supports continued U.S. leadership.



DynamoFL, Inc.
 2261 Market Street #4629
 San Francisco, CA 94114
 E: Christian@dynamo.ai

2. Dr. Lau, could you please describe the concept of technological singularity, share your perspective on whether you believe AI will achieve singularity, and provide an estimate on the possible timeframe for this development?

Artificial General Intelligence (AGI) is AI that can perform any cognitive task a human can, across domains, with general reasoning, learning, and adaptability. Many estimates expect leading general-purpose models to meet AGI definitions within the next decade. Technological singularity, on the other hand, refers to the point in time where AI systems are capable of self-improving recursively, independent from humans. This could theoretically get rid of the need for human intelligence to improve or advance these technologies. The timeline and likelihood of this development depend on many hard-to-predict factors such as the future of compute and energy costs, as well as the future autonomy and governance of AI systems. At this stage, it is unclear if AI will achieve singularity.

AGI represents a general-intelligence threshold beyond which a singularity becomes technically possible, but the singularity itself requires further leaps. As powerful AI systems evolve, so too will novel risks, requiring equally innovative technology and policy to promote human security, safety, and progress. It is essential that institutions adopt governance controls before the arrival of AGI, when the technology remains easier to supervise, evaluate, and control.

From Rep. Downing:

1. We are increasingly hearing that companies are using AI models from China and integrating them into the U.S. economy. How can we best position our financial markets to continue to innovate with AI while also protecting against potential risks from Chinese technologies?

The growth of Chinese-made AI Models and rise of Chinese Communist Party-influenced infiltration attacks on American-made large language models (LLMs) threaten U.S. sovereignty, transforming consumer AI technologies into attack vectors for infiltration and disruption of U.S. AI systems, mass collection and exploitation of American data, and targeted political influence operations against American institutions and citizens.

The rise of DeepSeek and other Chinese LLMs presents an emerging national security threat. Chinese-made models expose American institutions and citizens to: security risks, as CCP backdoors heighten model integrity concerns; privacy risks through mass data exploitation from access to large, sensitive datasets; and governance risks through censorship, output manipulation, and narrative bias. Dynamo's risk evaluation report of DeepSeek-R1 found the following key stats: the model had a 96.7% failure rate at adaptive jailbreak attacks, a 81% rate of generating malicious code, and a 52.5% rate of eliciting toxic or harmful content. Use of this



DynamoFL, Inc.
2261 Market Street #4629
San Francisco, CA 94114
E: Christian@dynamo.ai

model in high-impact financial industries would expose institutions and citizens to many potential risks.

And yet we've recently seen that even American-made LLMs are vulnerable to CCP influence, disruption, and exfiltration operations. CCP state actors use supply chain contamination and training data poisoning to infiltrate American AI systems, create pathways for CCP influence within US models, and complicate efforts to safeguard against foreign influence. As a result, American-made LLMs and AI agents can be vehicles for foreign interference and attacks.

As a result of these risks, it is imperative that financial institutions and governments deploying AI (1) adopt innovative test and evaluation techniques to assess third party and in-house models for key content and security risks, (2) employ real-time guardrails to block biased content, protect sensitive data, and block data poisoning attacks, (3) invest in comprehensive monitoring, observability, and threat detection technologies to ensure organizations are prepared to deal with emerging threats, and (4) support a healthy, competitive AI startup and talent ecosystem to continue to provide innovation to the financial services industry.

Congress can further promote these and other "defensive AI" technologies by:

1. Creating regulatory sandboxes to explore AI use cases and best practices for risk management;
2. Incentivizing and supporting a growing independent AI evaluation ecosystem;
3. Supporting discussions on the future of model risk management; and
4. Calling on financial and national security agencies to develop comprehensive plans to defend against adversarial AI-powered attacks, especially those brought on by advancements in Agentic AI.

November 24, 2025

The Honorable French Hill
Chairman, Committee on Financial Services
Committee on Banking, Housing, and Urban Affairs
U.S. House of Representatives
Washington, DC 20515

Dear Chairman Hill,

Enclosed is my response to the question for the record received after my appearance before the House Committee on Financial Services hearing held on September 18, 2025, titled "Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness."

Sincerely,

Dr. David Cox
Vice President for AI Models, IBM Research

Response to Question for Dr. David Cox
Vice President for AI Models, IBM Research

September 18, 2025 Committee on Financial Services hearing titled
“Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses, and Competitiveness”

Question from Rep. Downing:

We are increasingly hearing that companies are using AI models from China and integrating them into the U.S. economy. How can we best position our financial markets to continue to innovate with AI while also protecting against potential risks from Chinese technologies?

Answer:

A key challenge is recognizing that AI models developed by firms in the People’s Republic of China are not always served from China. Many are deployed through third-country infrastructure or can run locally on consumer hardware. Policymakers should therefore distinguish between models trained or served from data centers operating in mainland China, and models created by China-based companies but hosted, distributed, or executed outside PRC jurisdiction. These two vectors carry very different risk profiles for financial markets.

Regardless of origin, to evaluate any AI system we have to understand how it works: where it was developed, who can access it, and what data and methods underlie its training. The core risk of any proprietary technology is the trust gap: without visibility into its development, users cannot fully assess reliability, security, or resilience. This is why IBM continues to emphasize transparency as foundational to trustworthy AI.

One way to help promote trustworthy AI is by recognizing the value and role of open source AI. When AI models are developed openly, markets benefit from greater clarity, auditability, and shared innovation. Open source ecosystems can strengthen American competitiveness and security by:

- Allowing developers to focus on domain-specific innovations rather than rebuilding foundational components;
- Reducing exploitation risks and improving overall system security;
- Lowering barriers to market entry and enabling smaller firms to compete on a more equal footing;
- Reducing the cost of deploying AI that boosts productivity, strengthens cybersecurity, and accelerates revenue growth; and
- Expanding the talent pipeline by broadening access to tools that cultivate AI skills at scale

From Rep. Rose:

1. Mr. Reisman, could you please describe the concept of technological singularity, share your perspective on whether you believe AI will achieve singularity, and provide an estimate on the possible timeframe for this development?

- Technological singularity describes a hypothetical scenario in which artificial intelligence enters a continuous and accelerating cycle of learning that leads to “superintelligence” that exceeds human cognitive capabilities.
- Experts express a wide range of views as to if, or when, this scenario will take place, and how we will know if we have reached it. Some have suggested it could emerge within the coming few decades, and others that it might only arrive much later, if at all.
- From the perspective of public policy, it may be less useful to focus on a single, hypothetical future moment when AI develops in such a manner that it outperforms all aspects of human intelligence, but instead identify specific domains and use cases that AI is likely to transform most rapidly, with respect to both opportunities and risks, and make thoughtful plans to prepare accordingly.
- The Centre for Information Policy Leadership (CIPL) has published a number of works identifying best practices for ensuring responsible and beneficial development and use of AI that we recommend for further reference, including our reports on [Building Accountable AI Programs: Mapping Emerging Best Practices to the CIPL Accountability Framework](#) and on [Privacy-Enhancing and Privacy-Preserving Technologies in AI: Enabling Data Use and Operationalizing Privacy by Design and Default](#).

QUESTIONS FOR THE RECORD FROM THE HONORABLE JOHN ROSE

DANIEL S. GORFINE

CEO, GATTACA HORIZONS LLC; ADJUNCT PROFESSOR OF LAW,
GEORGETOWN UNIVERSITY LAW CENTER; FORMER CHIEF INNOVATION OFFICER
AND DIRECTOR, LABCFTC AT THE U.S. COMMODITY FUTURES TRADING
COMMISSION (CFTC)

BEFORE THE U.S. HOUSE FINANCIAL SERVICES
SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL TECHNOLOGY, AND
ARTIFICIAL INTELLIGENCE

“Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses,
and Competitiveness”

October 23, 2025

From Rep. Rose:

1. Mr. Gorfine, as you know, CFTC registration is a complex process that involves extensive recordkeeping. In your view, does artificial intelligence have the potential to significantly reduce the time and cost associated with meeting CFTC registration and recordkeeping requirements?

There is little question that artificial intelligence will enhance regulatory compliance, including during a regulatory application, licensure, or chartering process. AI is well placed to help regulated entities or applicants compile relevant information required by regulators. It can further improve ongoing recordkeeping and reporting functions, as well as help regulators analyze shared information. While there may be initial costs to procure, adopt, and train staff to use new AI tools, such tools should help to rapidly reduce overall costs and improve efficiencies.

2. Mr. Gorfine, while I recognize the significant promise that artificial intelligence holds in helping companies comply with complex recordkeeping requirements, I am concerned about the potential for companies to attribute inaccurate or misleading regulatory filings to their AI systems as a default defense. How can we balance embracing the benefits of AI with ensuring clear accountability when it comes to false or misleading regulatory Submissions?

While AI holds substantial promise and already is helping regulated entities be more efficient and effective across a range of regulatory or compliance functions, it is also critical that regulated entities maintain full accountability for their use of such AI tools. Depending on the AI use case or application, we often speak of the importance of having a “human in the loop” to verify and confirm AI-driven decisioning. To this end, regulators should communicate to regulated entities that they remain accountable for the accuracy of regulatory submissions, and that mere use of AI as a tool is not a defense to that accountability.

QUESTIONS FOR THE RECORD FROM THE HONORABLE TROY DOWNING

DANIEL S. GORFINE

CEO, GATTACA HORIZONS LLC; ADJUNCT PROFESSOR OF LAW,
GEORGETOWN UNIVERSITY LAW CENTER; FORMER CHIEF INNOVATION OFFICER
AND DIRECTOR, LABCFTC AT THE U.S. COMMODITY FUTURES TRADING
COMMISSION (CFTC)

BEFORE THE U.S. HOUSE FINANCIAL SERVICES
SUBCOMMITTEE ON DIGITAL ASSETS, FINANCIAL TECHNOLOGY, AND
ARTIFICIAL INTELLIGENCE

“Unlocking the Next Generation of AI in the U.S. Financial System for Consumers, Businesses,
and Competitiveness”

October 23, 2025

From Rep. Downing:

1. Could you describe some of most concerning aspects of the Biden Administration's approach to AI, from executive orders to his so-called "AI Bill of Rights," and discuss what the impact would have been had Congress passed what the Biden Administration proposed?

As noted in my written testimony, financial services law, regulations, and risk management frameworks have long governed the adoption of emerging technologies in the sector, whether through rules to ensure consumer protection or principles-based guidance to ensure safety and soundness. These laws, regulations, and risk management principles largely apply to conduct and activities regardless of the technological tools used by the regulated entity—in this way they are appropriately technology-neutral and fit-for-purpose to govern AI adoption. Against this backdrop, the tone set by regulators can have a "soft power" impact on financial institutions looking to adopt AI. During the Biden Administration, a number of key financial regulators expressed strong skepticism of AI-driven approaches, likely deterring adoption, especially for smaller firms lacking large compliance teams capable of parsing regulatory expectations.

An example of a regulatory proposal that was overbroad and premature in its requirements related to AI-technologies was the SEC's proposed rulemaking regarding "Conflicts of Interest Associated with the Use of Predictive Data Analytics by Broker Dealers and Investment Advisers" during the past Biden Administration. The proposed rule offered an overbroad definition of covered technologies, including AI technologies that have long been in existence, implied and asserted that such technologies pose a greater risk of harm to investors and going undetected (despite clear historical evidence of misconduct made harder to trace by legacy tools, such as unscrupulous brokers targeting investors by telephone), and then sought to treat firms using these technologies differently from those relying on legacy systems by subjecting them to a more punitive conflict of interest framework.

By specifically targeting emerging technologies largely based on conduct already subject to regulation or speculative conduct that may (or may not) occur in the future, the proposal was inherently not technology-neutral and would have deterred adoption of such technologies, especially for smaller firms that lack large compliance teams capable of parsing ambiguous regulatory expectations. Fortunately, this proposed rule has been withdrawn.

Questions for the Record**Subcommittee on Digital Assets, Financial Technology, and Artificial Intelligence of the
Committee on Financial Services****Hearing, titled: Unlocking the Next Generation of AI in the U.S. Financial System for
Consumers, Businesses, and Competitiveness****Responses from Nicol Turner Lee, Brookings Institution****October 23, 2025**From Rep. Pressley:

1. Dr. Turner Lee, thank you for your answers about how a national financial literacy program might be helpful to include more people in shaping AI in financial services and center their lived experiences so that financial tools better suit them and address economic disparities. I want to follow up on my question about how the federal government can work to ensure that people from all walks of life can have a role in shaping and working in AI, especially in the financial services sector. How can companies also lead in efforts to include all people in shaping and working in AI within financial services?

Congress can facilitate and support a national financial literacy campaign to elevate the awareness of communities and individuals about how artificial intelligence (AI) is being used in the financial services sector. The campaign can particularly target information and strategies for averting scams, fraudulent activities, and other malicious behaviors of bad actors. Currently, seniors are exposed to financial abuse with many finding themselves exposed to egregious romance scams and other manipulative tactics to extract their wealth. According to a report by the Federal Trade Commission (FTC), between 2020 and 2024, there was a four-fold increase in reports of impersonation scams from seniors who lost more than \$10,000.¹ Additionally, there was an eight-fold increase in reports from seniors who lost more than \$100,000.² In total, scams targeting adults over 60 caused over \$3.4 billion in losses in 2023.³ Deloitte's Center for

¹ Federal Trade Commission. "FTC Data Show a More Than Four-Fold Increase in Reports of Impersonation Scammers Stealing Tens and Even Hundreds of Thousands from Older Adults," August 7, 2025, <https://www.ftc.gov/news-events/news/press-releases/2025/08/ftc-data-show-more-four-fold-increase-reports-impersonation-scammers-stealing-tens-even-hundreds>.

² Ibid.

³ Federal Bureau of Investigation Internet Crime Complaint Center. "Elder Fraud, In Focus," April 30, 2024, <https://www.fbi.gov/news/stories/elder-fraud-in-focus>.

Financial Services predicts that generative AI could enable fraud losses to reach \$40 billion in the U.S. by 2027, up from \$12.3 billion in 2023.⁴

To address these and other issues related to AI's use in consumer financial services, it is also important for Congress to provide information to consumers who may or may not have internet access. For example, generative AI can extract and mimic the voices, images, or dialogue of trusted loved ones, who—through a simple phone call or text message—can impersonate the voice or mirror the writing style of a loved one asking for monetary support. Congress has not adequately addressed the role of deepfakes as AI technology has become more sophisticated. This leaves open more opportunities for bad actors—whether in the U.S. or abroad—to attack.

Further, financial services companies should be part of a national financial literacy campaign to advance AI that protects consumers and provides assurances around the trustworthiness of their applications. In my written and oral testimonies, I expressed concerns about the use of regulatory sandboxes that do not include consumers, which is why there should be transparency and congressional oversight as these are designed and implemented. AI companies without guardrails present challenges for the preservation of consumer protections. Additionally, consumers subjected to harm have little to no recourse to mitigate risk. Regulatory sandboxes and safe harbors must be done in partnership with communities and consumers. Better yet, they should be perceived more like “greenhouses” with open transparency around the goals of the project. Other countries, such as Singapore, offer a model for how sandboxes that prioritize fairness and accountability could work.⁵ This includes rigorously testing innovations to identify and address gaps and gathering insights from diverse groups of stakeholders.⁶

In sum, Congress can and should support a national financial literacy program centered on educating consumers about AI use in this sector and raising awareness about the growing influence of scams and frauds enabled by AI. Organizations that support financial literacy and counseling should be actively engaged in such national awareness efforts. If Congress had a clear and enforceable national data privacy standard that addresses third-party use of data, some of the challenges of financial misinformation could be adequately addressed, at least for now.

2. Dr. Turner Lee, you also mentioned that you think there are ways the government can incentivize the private sector to be more inclusive of diverse founders, entrepreneurs, and small businesses who want to participate in AI work within financial services. What are some of the incentives the government should use to encourage the private sector to be more inclusive?

⁴ Satish Lalchand et al. “Generative AI Is Expected to Magnify the Risk of Deepfakes and Other Fraud in Banking,” *Deloitte Insights*, May 29, 2024, <https://www.deloitte.com/us/en/insights/industry/financial-services/deepfake-banking-fraud-risk-on-the-rise.html>.

⁵ Richard Sentinella. “How different jurisdictions approach AI regulatory sandboxes,” IAPP, May 14, 2025. <https://iapp.org/news/a/how-different-jurisdictions-approach-ai-regulatory-sandboxes>.

⁶ *Ibid.*

The United States continues to have a gap in capital investments in small businesses, especially those created and led by diverse founders. In 2022, Black and Latino founders received just 1-1.5% of venture capital funding among all startups in the U.S.⁷ Investments in startups founded by Black women are even more daunting; in 2021, Black women received just 0.34% of venture capital funding.⁸ Leveling the playing field and supporting more diverse founders will require some support from Congress and the agencies tasked with cultivating small businesses. Policies that ensure everyone has access to explore their interests, attain relevant education, and be exposed to a variety of opportunities can inspire the next generation of diverse founders.

Congress can also task the Small Business Administration (SBA) with a review of the obstacles to capital acquisition by diverse founders to inform their work and focus on scaled education for entrepreneurs and innovators. In the current environment, the rollbacks on diversity, equity, and inclusion have made such efforts more difficult, which ultimately impacts the economic growth of U.S. small businesses of all types with all ranges of founders to lead the future of economic prosperity. Supporting policies that enable startups and innovators to also transact with Global South countries and regions might also establish new markets for diverse startups and entrepreneurs.

3. Dr. Turner Lee, what should a national financial literacy program be sure to include as we pursue both the ethical and innovative use of AI that includes all?

The messaging that should be part of a national financial literacy program should answer these questions:

- What is AI and how is it being used in financial services?
- What are ways that AI can help with certain financial functions? (e.g., budgeting, savings, etc.)
- What agency do users have when AI is used in financial decisions? How can they take action?
- How are fraud and scams (e.g., video, text messaging, audio) being enabled by AI? What do users need to know to avoid being victimized by scams and fraudulent activities?

⁷ Ganesan et al. "Underestimated Start-up Founders: The Untapped Opportunity," McKinsey and Company, June 23, 2023, <https://www.mckinsey.com/featured-insights/diversity-and-inclusion/underestimated-start-up-founders-the-untapped-opportunity>.

⁸ Sophia Kunthara. "Black Women Still Receive Just a Tiny Fraction of VC Funding Despite 5-Year High," Crunchbase News, July 16, 2021, <https://news.crunchbase.com/diversity/something-ventured-black-women-founders/>.

- What checklist should online users reference to ensure that online users are not falling prey to financial manipulation?
- Who should online users contact if they feel negatively impacted by AI?

Clearly, AI could help democratize access to financial planning and wealth building to help narrow the growing income divide and close the racial wealth gap. However, AI technology is imperfect and when in the hands of bad actors can impact personal identity, particularly as it relates to one's financial security. Further, for economically marginalized populations, any action that negatively impacts their wealth can foreclose future opportunities, deepening the existing wealth gap. Literacy programs could be offered in schools and libraries while also made available by banks and other financial institutions. Companies with models that supply financial information should also be required to have an independent audit and make those results publicly available. Just as compliance rules exist for financial advisers, similar requirements should be met by models who espouse the same capabilities as trusted professionals. Companies should clearly disclose how and when these models should be used to help people better discern what types of decisions could be trusted with AI. Having Congress assess and enact measures that promote public disclosure of AI use in financial products and decision-making will reduce the potential harms that impact online users and make AI use less opaque. Congress should also actively develop legislation that brings more safeguards to consumers to repel the intrusions of third-party companies that are regularly leveraging AI for scams and frauds.

119TH CONGRESS
1ST SESSION

H. R. 4801

To establish AI Innovation Labs that permit certain persons to experiment with artificial intelligence without expectation of enforcement actions.

IN THE HOUSE OF REPRESENTATIVES

JULY 29, 2025

Mr. HILL of Arkansas (for himself, Mr. TORRES of New York, Mr. STEIL, and Mr. GOTTHEIMER) introduced the following bill; which was referred to the Committee on Financial Services

A BILL

To establish AI Innovation Labs that permit certain persons to experiment with artificial intelligence without expectation of enforcement actions.

1 *Be it enacted by the Senate and House of Representa-*
2 *tives of the United States of America in Congress assembled,*

3 **SECTION 1. SHORT TITLE.**

4 This Act may be cited as the “Unleashing AI Innova-
5 tion in Financial Services Act”.

6 **SEC. 2. DEFINITIONS.**

7 In this section:

8 (1) AI TEST PROJECT.—The term “AI test
9 project” means a financial product or service that—

1 (A) falls under the jurisdiction of a finan-
2 cial regulatory agency;

3 (B) makes substantial use of artificial in-
4 telligence; and

5 (C) is, or may be, subject to a Federal reg-
6 ulation or Federal statute.

7 (2) APPROPRIATE FINANCIAL REGULATORY
8 AGENCY.—The term “appropriate financial regu-
9 latory agency” means—

10 (A) the appropriate Federal banking agen-
11 cy, as defined in section 3 of the Federal De-
12 posit Insurance Act (12 U.S.C. 1813), with re-
13 spect to an institution described in subsection
14 (q) of that section;

15 (B) the Securities and Exchange Commis-
16 sion, with respect to an institution not de-
17 scribed in subparagraph (A) that is—

18 (i) any broker or dealer that is reg-
19 istered with the Commission under the Se-
20 curities Exchange Act of 1934 (15 U.S.C.
21 78a et seq.);

22 (ii) any investment company that is
23 registered with the Commission under the
24 Investment Company Act of 1940 (15
25 U.S.C. 80a–1 et seq.);

1 (iii) any investment adviser that is
2 registered with the Commission under the
3 Investment Advisers Act of 1940 (15
4 U.S.C. 80b-1 et seq.);

5 (iv) any clearing agency registered
6 with the Commission under the Securities
7 Exchange Act of 1934 (15 U.S.C. 78a et
8 seq.);

9 (v) any nationally recognized statis-
10 tical rating organization registered with
11 the Commission under the Securities Ex-
12 change Act of 1934 (15 U.S.C. 78a et
13 seq.);

14 (vi) any transfer agent registered with
15 the Commission under the Securities Ex-
16 change Act of 1934 (15 U.S.C. 78a et
17 seq.);

18 (vii) any exchange registered as a na-
19 tional securities exchange with the Com-
20 mission under the Securities Exchange Act
21 of 1934 (15 U.S.C. 78a et seq.);

22 (viii) any national securities associa-
23 tion registered with the Commission under
24 the Securities Exchange Act of 1934 (15
25 U.S.C. 78a et seq.);

1 (ix) any securities information proc-
2 essor registered with the Commission
3 under the Securities Exchange Act of 1934
4 (15 U.S.C. 78a et seq.);

5 (x) the Municipal Securities Rule-
6 making Board established under the Secu-
7 rities Exchange Act of 1934 (15 U.S.C.
8 78a et seq.);

9 (xi) the Public Company Accounting
10 Oversight Board established under the
11 Sarbanes-Oxley Act of 2002 (15 U.S.C.
12 7211 et seq.);

13 (xii) the Securities Investor Protection
14 Corporation established under the Securi-
15 ties Investor Protection Act of 1970 (15
16 U.S.C. 78aaa et seq.); and

17 (xiii) any security-based swap execu-
18 tion facility, security-based swap data re-
19 pository, security-based swap dealer, or
20 major security-based swap participant reg-
21 istered with the Commission under the Se-
22 curities Exchange Act of 1934 (15 U.S.C.
23 78a et seq.), with respect to the security-
24 based swap activities of the person that re-

1 quire such person to be registered under
2 such Act;

3 (C) the Bureau of Consumer Financial
4 Protection, with respect to a covered person, as
5 defined in section 1002 of the Consumer Finan-
6 cial Protection Act of 2010 (12 U.S.C. 5481),
7 that does not have an appropriate financial reg-
8 ulatory agency under subparagraph (A), (B),
9 (D), or (E) of this paragraph;

10 (D) the National Credit Union Administra-
11 tion, with respect to an insured credit union, as
12 defined in section 101 of the Federal Credit
13 Union Act (12 U.S.C. 1752); and

14 (E) the Federal Housing Finance Agency,
15 with respect to—

16 (i) a Federal Home Loan Bank;

17 (ii) the Federal Home Loan Bank
18 System;

19 (iii) the Federal National Mortgage
20 Association; and

21 (iv) the Federal Home Loan Mortgage
22 Corporation.

23 (3) ARTIFICIAL INTELLIGENCE; AI.—The terms
24 “artificial intelligence” and “AI” have the meaning
25 given the term “artificial intelligence” in section

1 5002 of the National Artificial Intelligence Initiative
2 Act of 2020 (15 U.S.C. 9401).

3 (4) COMMISSION.—The term “Commission”
4 means the Securities and Exchange Commission.

5 (5) FEDERAL SECURITIES LAWS.—The term
6 “Federal securities laws” means—

7 (A) the Securities Act of 1933 (15 U.S.C.
8 77a et seq.);

9 (B) the Securities Exchange Act of 1934
10 (15 U.S.C. 78a et seq.);

11 (C) the Sarbanes-Oxley Act of 2002 (15
12 U.S.C. 7201 et seq.);

13 (D) the Trust Indenture Act of 1939 (15
14 U.S.C. 77aaa et seq.);

15 (E) the Investment Company Act of 1940
16 (15 U.S.C. 80a–1 et seq.);

17 (F) the Investment Advisers Act of 1940
18 (15 U.S.C. 80b–1 et seq.);

19 (G) the Jumpstart Our Business Startup
20 Act (Public Law 112–106; 126 Stat. 306); and

21 (H) the Dodd-Frank Wall Street Reform
22 and Consumer Protection Act (Public Law
23 111–203; 124 Stat. 1376).

24 (6) FINANCIAL PRODUCT OR SERVICE.—The
25 term “financial product or service”—

1 (A) has the meaning given the term in sec-
2 tion 1002 of the Consumer Financial Protection
3 Act of 2010 (12 U.S.C. 5481);

4 (B) includes—

5 (i) activities that are financial in na-
6 ture, as defined in section 4(k)(4) of the
7 Bank Holding Company Act of 1956 (12
8 U.S.C. 1843(k)(4));

9 (ii) any financial product or service
10 provided by a person regulated by the
11 Commission, as defined in section 1002 of
12 the Consumer Financial Protection Act of
13 2010 (12 U.S.C. 5481); and

14 (iii) includes the offer or sale of any
15 security subject to the Federal securities
16 laws; and

17 (C) does not include the business of insur-
18 ance.

19 (7) FINANCIAL REGULATORY AGENCY.—The
20 term “financial regulatory agency” means—

21 (A) the Board of Governors of the Federal
22 Reserve System;

23 (B) the Federal Deposit Insurance Cor-
24 poration;

1 (C) the Office of the Comptroller of the
2 Currency;

3 (D) the Securities and Exchange Commis-
4 sion;

5 (E) the Bureau of Consumer Financial
6 Protection;

7 (F) the National Credit Union Administra-
8 tion; and

9 (G) the Federal Housing Finance Agency.

10 (8) REGULATED ENTITY.—The term “regulated
11 entity” means an entity regulated by any financial
12 regulatory agency.

13 **SEC. 3. USE OF ARTIFICIAL INTELLIGENCE BY REGULATED**
14 **FINANCIAL ENTITIES.**

15 (a) AI INNOVATION LABS.—

16 (1) ESTABLISHMENT.—Each financial regu-
17 latory agency shall establish, or identify an office,
18 division, or department of the agency that shall
19 serve as, an AI Innovation Lab to enable regulated
20 entities to experiment with AI test projects without
21 unnecessary or unduly burdensome regulation or ex-
22 pectation of enforcement actions, pursuant to the
23 approval of an application under paragraph (2).

24 (2) APPLICATIONS.—

25 (A) SUBMISSION.—

1 (i) IN GENERAL.—A regulated entity
2 may submit to the appropriate financial
3 regulatory agency an application, on a
4 form determined by the appropriate finan-
5 cial regulatory agency, to engage in an AI
6 test project through the AI Innovation Lab
7 established or identified under paragraph
8 (1).

9 (ii) CONTENTS.—An application sub-
10 mitted under clause (i) shall include—

11 (I) a description of the AI test
12 project proposed to be carried out by
13 the regulated entity;

14 (II) an alternative compliance
15 strategy that—

16 (aa) identifies a regulation
17 issued by the appropriate finan-
18 cial regulatory agency that the
19 regulated entity requests to be
20 waived or modified; and

21 (bb) proposes an alternative
22 method for the regulated entity
23 to comply with the regulation, in-
24 cluding an explanation as to why
25 the alternative method is essen-

1 tial to the operation of the entity
2 and how the regulated entity
3 would effectively manage risks
4 associated with the AI test
5 project;

6 (III) an explanation of how under
7 the strategy described in subclause
8 (II), the AI test project—

9 (aa) would serve the public
10 interest, improve consumer or in-
11 vestor access to a financial prod-
12 uct or service, or promote con-
13 sumer or investor protection;

14 (bb) would enhance effi-
15 ciency or operations, foster inno-
16 vation or competitiveness, im-
17 prove risk management and secu-
18 rity, or enhance regulatory com-
19 pliance;

20 (cc) would not present a sys-
21 temic risk to the financial system
22 of the United States;

23 (dd) is consistent with the
24 purposes of the anti-money laun-
25 dering and countering the financ-

1 ing of terrorism obligations under
2 subchapter II of chapter 53 of
3 title 31, United States Code; and
4 (ee) would not present a na-
5 tional security risk to the United
6 States;

7 (IV) a proposed date on which
8 the AI test project would terminate
9 and an explanation why such termi-
10 nation date would be appropriate;

11 (V) proposed limitations on the
12 size, scope, and growth of the AI test
13 project;

14 (VI) a detailed business plan;
15 and

16 (VII) an estimate of the eco-
17 nomic impact of the AI test project if
18 approved.

19 (iii) JOINT APPLICATIONS.—Two or
20 more regulated entities may submit a joint
21 application to the same financial regu-
22 latory agency under clause (i).

23 (iv) REGULATIONS OF OTHER AGEN-
24 CIES.—

1 (I) IN GENERAL.—A regulated
2 entity may submit an application
3 under this subparagraph that includes
4 an alternative compliance strategy for
5 a regulation issued or enforced by a
6 financial regulatory agency that is not
7 the appropriate financial regulatory
8 agency for the regulated entity.

9 (II) REQUIREMENTS.—An appli-
10 cation described in subclause (I) shall
11 be subject to the same requirements
12 as an application described in clause
13 (ii), except that—

14 (aa) the regulated entity
15 shall submit the application to
16 the appropriate financial regu-
17 latory agency and the financial
18 regulatory agency that issued or
19 enforces the regulation that is
20 the subject of the alternative
21 compliance strategy; and

22 (bb) the AI test project may
23 not take effect unless the appro-
24 priate financial regulatory agency
25 and any other financial regu-

1 latory agency that issued or en-
2 forces the regulation that is the
3 subject of the alternative compli-
4 ance strategy jointly approve the
5 application using the process de-
6 scribed in subparagraph (B).

7 (v) NOTICE.—A regulated entity that
8 is regulated or supervised by more than 1
9 financial regulatory agency shall provide
10 notice of any application submitted to the
11 appropriate financial regulatory agency
12 under this section to each financial regu-
13 latory agency by which it is regulated or
14 supervised not later than 5 business days
15 after the entity submits the application to
16 the appropriate financial regulatory agen-
17 cy.

18 (B) AGENCY REVIEW.—

19 (i) IN GENERAL.—Except as provided
20 in clause (iv), not later than 120 days
21 after the date on which an application is
22 submitted to the appropriate financial reg-
23 ulatory agency under subparagraph (A),
24 the appropriate financial regulatory agency
25 shall—

- 1 (I) review the application; and
2 (II) submit to the applicant in
3 writing a determination of the agency.

4 (ii) APPROVAL.—

- 5 (I) IN GENERAL.—If the appli-
6 cant shows that it is more likely than
7 not that the application meets the re-
8 quirements for establishing an alter-
9 native compliance strategy and satis-
10 fies the standards described in sub-
11 clauses (II) and (III) of subparagraph
12 (A)(ii), the agency shall approve the
13 application and notify the applicant in
14 writing of—

- 15 (aa) the regulation that is
16 the subject of the alternative
17 compliance strategy;

- 18 (bb) the terms of the alter-
19 native compliance strategy for
20 the AI test project;

- 21 (cc) the date on which the
22 AI test project will terminate;

- 23 (dd) any limitations on the
24 size, scope, or growth of the AI
25 test project; and

1 (ee) any additional limita-
2 tions or conditions on the AI test
3 project, as determined by the ap-
4 propriate financial regulatory
5 agency.

6 (II) EFFECT OF APPROVAL.—
7 With respect to an AI test project, ex-
8 cept as provided in subclause (III),
9 beginning on the date on which an ap-
10 plication submitted under subpara-
11 graph (A) is approved and ending on
12 the date described in subclause
13 (I)(cc)—

14 (aa) the appropriate finan-
15 cial regulatory agency may en-
16 force a regulation described in
17 subclause (I)(aa) only in the
18 manner set out in the alternative
19 compliance strategy described in
20 subclause (I)(bb); and

21 (bb) a financial regulatory
22 agency that is not the appro-
23 priate financial regulatory agency
24 may not enforce a regulation de-
25 scribed in subclause (I)(aa).

1 (III) ENFORCEMENT BY AN-
2 OTHER FINANCIAL REGULATORY
3 AGENCY.—With respect to an AI test
4 project, a financial regulatory agency
5 other than the appropriate financial
6 regulatory agency that approves an
7 application under subparagraph
8 (A)(iv) may enforce a regulation de-
9 scribed in subclause (I)(aa) if the al-
10 ternative compliance strategy de-
11 scribed in subclause (I)(bb) provides
12 for enforcement by such financial reg-
13 ulatory agency.

14 (IV) RULE OF CONSTRUCTION.—
15 Nothing in this clause may be con-
16 strued to limit the authority of a fi-
17 nancial regulatory agency to take an
18 enforcement action against a regu-
19 lated entity with respect to fraud or
20 for engaging in an unsafe or unsound
21 practice relating to an AI test project.

22 (iii) DENIAL.—

23 (I) IN GENERAL.—If an agency
24 denies an application submitted under
25 subparagraph (A), the agency—

1 (aa) shall submit to the ap-
2 plicant a written notice explain-
3 ing the reason for denial; and

4 (bb) may not take an en-
5 forcement action related to the
6 proposed AI test project against
7 the applicant earlier than the
8 date that is 30 days after the
9 date on which the agency submits
10 the written notice described in
11 item (aa).

12 (II) RESUBMITTALS.—Each time
13 an application submitted under sub-
14 paragraph (A) is denied, the regulated
15 entity—

16 (aa) may submit an amend-
17 ed application after receiving
18 feedback from the agency making
19 such denial; and

20 (bb) may not resubmit more
21 than 2 applications that are sub-
22 stantially similar to the denied
23 application.

24 (III) INJUNCTIVE RELIEF.—A fi-
25 nancial regulatory agency, by and

1 through its own attorneys, may file a
2 civil action in an appropriate United
3 States district court to enjoin an AI
4 test project if the agency determines
5 that the AI test project—

6 (aa) presents an immediate
7 danger to consumers or investors;
8 or

9 (bb) presents a risk—

10 (AA) to financial mar-
11 kets;

12 (BB) in the case of an
13 AI test project engaged in
14 by an insured depository in-
15 stitution or an insured credit
16 union, of loss to a Federal
17 deposit or share insurance
18 fund;

19 (CC) of a violation of
20 anti-money laundering and
21 countering the financing of
22 terrorism obligations under
23 subchapter II of chapter 53
24 of title 31, United States
25 Code; or

1 (DD) to the national
2 security of the United
3 States.

4 (iv) EXTENSION.—If the financial reg-
5 ulatory agency needs additional time, the
6 agency may extend the approval deadline
7 by 120 days. After the expiration of the
8 120-day extension period, if the agency has
9 not made a determination on the applica-
10 tion, the application will automatically be
11 deemed approved and effective.

12 (C) DATA SECURITY.—All data supplied by
13 sponsors of AI test projects to a financial regu-
14 latory agency submitted under this section shall
15 be stored and maintained in a secure manner
16 by the financial regulatory agency, consistent
17 with applicable data security standards.

18 (D) REGULATIONS.—Not later than 180
19 days after the date of enactment of this Act,
20 each financial regulatory agency shall promul-
21 gate regulations that—

22 (i) shall be published in the Federal
23 Register and provide a 60-day period for
24 public notice and comment; and

25 (ii) include—

1 (I) procedures for modifying the
2 AI test projects that are approved by
3 the agency;

4 (II) consequences for failure to
5 comply with the terms of an alter-
6 native compliance strategy;

7 (III) a requirement that an AI
8 test project will terminate not earlier
9 than 1 year after the AI test project
10 is approved;

11 (IV) procedures to extend the
12 termination date described in sub-
13 clause (III);

14 (V) procedures for confiden-
15 tiality; and

16 (VI) procedures for coordinating
17 decisions relating to applications sub-
18 mitted jointly by multiple regulated
19 entities or applications submitted to
20 more than one financial regulatory
21 agency.

22 (b) REPORT.—Not later than 2 years after the date
23 of enactment of this Act, and each year for 7 years there-
24 after, each financial regulatory agency shall submit to the
25 Committee on Banking, Housing, and Urban Affairs of

1 the Senate and the Committee on Financial Services of
2 the House of Representatives an annual report on the out-
3 comes of AI test projects. A report under this subsection
4 may not include the names of participating entities or any
5 proprietary or confidential business information. A report
6 under this subsection shall include aggregated findings,
7 trends, and lessons learned from the AI test projects.

8 (c) RULE OF CONSTRUCTION.—Nothing in this sec-
9 tion may be construed to limit the authority of a financial
10 regulatory agency to take an enforcement action against
11 a regulated entity with respect to fraud relating to an AI
12 test project.

○

119TH CONGRESS
1ST SESSION

H. R. 2152

To require a strategy to defend against the economic and national security risks posed by the use of artificial intelligence in the commission of financial crimes, including fraud and the dissemination of misinformation, and for other purposes.

IN THE HOUSE OF REPRESENTATIVES

MARCH 14, 2025

Mr. NUNN of Iowa (for himself and Mr. HIMES) introduced the following bill;
which was referred to the Committee on Financial Services

A BILL

To require a strategy to defend against the economic and national security risks posed by the use of artificial intelligence in the commission of financial crimes, including fraud and the dissemination of misinformation, and for other purposes.

1 *Be it enacted by the Senate and House of Representa-*
2 *tives of the United States of America in Congress assembled,*

3 **SECTION 1. SHORT TITLE.**

4 This Act may be cited as the “Artificial Intelligence
5 Practices, Logistics, Actions, and Necessities Act” or the
6 “AI PLAN Act”.

1 **SEC. 2. STRATEGY TO DEFEND AGAINST RISKS POSED BY**
2 **THE USE OF ARTIFICIAL INTELLIGENCE.**

3 (a) SENSE OF CONGRESS.—It is the sense of Con-
4 gress that the development and use of artificial intelligence
5 in the commission of financial crimes by adversarial actors
6 poses a significant risk to the national and economic secu-
7 rity of the United States.

8 (b) STRATEGY TO DEFEND AGAINST RISKS POSED
9 BY MISINFORMATION, FRAUD, AND FINANCIAL CRIME
10 CONDUCTED WITH ARTIFICIAL INTELLIGENCE.—

11 (1) IN GENERAL.—Not later than 180 days
12 after the date of the enactment of this Act and an-
13 nually thereafter, the Secretary of the Treasury, the
14 Secretary of Homeland Security, and the Secretary
15 of Commerce, in consultation with the officials speci-
16 fied in paragraph (3), shall jointly submit to Con-
17 gress a report that includes the following:

18 (A) A description of interagency policies
19 and procedures to defend United States finan-
20 cial markets, United States persons, United
21 States businesses, and global supply chains
22 from the national and economic security risks
23 posed by the use of artificial intelligence in the
24 commission of financial crimes, including fraud
25 and the dissemination of misinformation.

1 (B) An itemized list of readily available re-
2 sources, hardware, software, and technologies
3 that can be immediately utilized to combat the
4 use of artificial intelligence in the commission
5 of financial crimes, including fraud and the dis-
6 semination of misinformation.

7 (C) An itemized list of resources, hard-
8 ware, software, technologies, people, and budg-
9 etary estimates needed to help Federal depart-
10 ments and agencies to combat the use of artifi-
11 cial intelligence in the commission of financial
12 crimes, including fraud and the dissemination
13 of misinformation.

14 (2) CONSIDERATIONS.—Reports required pur-
15 suant to paragraph (1) shall take the following risks
16 into consideration:

17 (A) Deepfakes.

18 (B) Voice cloning.

19 (C) Foreign election interference.

20 (D) Synthetic Identities.

21 (E) False flags and false signals that dis-
22 rupt market operations.

23 (F) Overall digital fraud.

24 (3) OFFICIALS SPECIFIED.—The officials speci-
25 fied in this paragraph are the following:

1 (A) The United States Trade Representa-
2 tive.

3 (B) The Attorney General.

4 (C) The Chairman of the Board of Gov-
5 ernors of the Federal Reserve System.

6 (D) The Director of the National Institute
7 of Standards and Technology.

8 (E) The Under Secretary of Commerce for
9 Industry and Security.

10 (F) The Chairman of the Securities and
11 Exchange Commission.

12 (c) RECOMMENDATIONS.—Not later than 90 days
13 after each report under subsection (b) is submitted, the
14 Secretary of the Treasury, the Secretary of Homeland Se-
15 curity, and the Secretary of Commerce shall jointly submit
16 to Congress a set of recommendations relating to each
17 such respective report that contain the following:

18 (1) Legislative recommendations to address the
19 risks posed by the use of artificial intelligence in the
20 commission of financial crimes, including fraud and
21 the dissemination of misinformation.

22 (2) Best practices to assist American businesses
23 and government entities with risk mitigation and in-
24 cident response to address the risks posed by the use
25 of artificial intelligence in the commission of finan-

1 cial crimes, including fraud and the dissemination of
2 misinformation.

○

119TH CONGRESS
1ST SESSION

H. R. 1734

To establish the Task Force on Artificial Intelligence in the Financial Services Sector to report to Congress on issues related to artificial intelligence in the financial services sector, and for other purposes.

IN THE HOUSE OF REPRESENTATIVES

FEBRUARY 27, 2025

Ms. PETTERSEN (for herself, Mr. FLOOD, Ms. BONAMICI, Mr. FITZPATRICK, Ms. OCASIO-CORTEZ, and Mr. HARDER of California) introduced the following bill; which was referred to the Committee on Financial Services

A BILL

To establish the Task Force on Artificial Intelligence in the Financial Services Sector to report to Congress on issues related to artificial intelligence in the financial services sector, and for other purposes.

1 *Be it enacted by the Senate and House of Representa-*
2 *tives of the United States of America in Congress assembled,*

3 **SECTION 1. SHORT TITLE.**

4 This Act may be cited as the “Preventing Deep Fake
5 Scams Act”.

6 **SEC. 2. FINDINGS.**

7 The Congress finds the following:

1 (1) Artificial intelligence is being used in new
2 and innovative ways by the financial services sector.

3 (2) Artificial intelligence may provide benefits
4 to banks, credit unions, and banking consumers.

5 (3) Artificial intelligence poses unique threats
6 to the safety and security of customer accounts.

7 (4) Voice banking is offered by many banks for
8 security and convenience reasons.

9 (5) The popularity of social media has made
10 video and audio of potential targets easier to obtain
11 for bad actors. These materials can be exploited to
12 replicate the voices and appearances of other people
13 in pursuit of data theft, identity theft, or fraud.

14 (6) Bad actors could utilize “deep fakes”, in-
15 cluding voice and audio manipulation, to compromise
16 and access a consumer’s financial accounts.

17 **SEC. 3. TASK FORCE ON ARTIFICIAL INTELLIGENCE IN THE**
18 **FINANCIAL SERVICES SECTOR.**

19 (a) ESTABLISHMENT.—There is established a Task
20 Force on Artificial Intelligence in the Financial Services
21 Sector (in this section referred to as the “Task Force”).

22 (b) MEMBERSHIP.—The Task Force shall consist of
23 the following:

24 (1) The Secretary of the Treasury, or a des-
25 ignee, who shall serve as Chair of the Task Force.

1 (2) The Comptroller of the Currency, or a des-
2 ignee.

3 (3) The Chairman of the Board of Governors of
4 the Federal Reserve System, or a designee.

5 (4) The Chairperson of the Federal Deposit In-
6 surance Corporation, or a designee.

7 (5) The Director of the Bureau of Consumer
8 Financial Protection, or a designee.

9 (6) The Chairman of the National Credit Union
10 Administration, or a designee.

11 (7) The Director of the Financial Crimes En-
12 forcement Network, or a designee.

13 (c) REPORT.—

14 (1) IN GENERAL.—Not later than the end of
15 the 1-year period beginning on the date of enact-
16 ment of this Act, the Task Force shall issue a report
17 to Congress containing the contents described in
18 paragraph (3).

19 (2) CONSULTATION.—

20 (A) REQUEST FOR INFORMATION.—Not
21 later than 90 days after the date of enactment
22 of this Act, the Task Force shall solicit public
23 feedback on the report required under para-
24 graph (1).

1 (B) INDUSTRY AND EXPERT STAKE-
2 HOLDERS.—In developing the report required
3 under paragraph (1), the Task Force shall seek
4 out and consult with industry and expert stake-
5 holders, including—

6 (i) depository institutions of varying
7 asset sizes;

8 (ii) credit unions of varying asset
9 sizes;

10 (iii) third-party vendors who use arti-
11 ficial intelligence when providing services
12 to depository institutions and credit
13 unions; and

14 (iv) artificial intelligence experts.

15 (3) CONTENTS.—The contents of the report are
16 the following:

17 (A) A description of how banks and credit
18 unions proactively protect themselves and con-
19 sumers from fraud utilizing artificial intel-
20 ligence.

21 (B) A list of standard definitions for the
22 different manners in which artificial intelligence
23 is used, including terms like “generative AI”,
24 “machine learning”, “natural language proc-
25 essing”, “algorithmic AI”, and “deep fakes”.

1 (C) A description of potential risks that
2 could result from the use of artificial intel-
3 ligence by bad actors to steal consumers' data,
4 steal consumers' identities, and commit fraud.

5 (D) A list of best practices for financial in-
6 stitutions to protect their customers from at-
7 tempts to steal consumer data, steal consumers'
8 identities, or commit fraud.

9 (E) Legislative and regulatory rec-
10 ommendations for the regulation of artificial in-
11 telligence and to protect consumers from data
12 theft, identity theft, and fraud.

13 (d) TERMINATION.—The Task Force shall terminate
14 90 days after the final report is issued pursuant to sub-
15 section (c).

○