

The Honorable John Thune
Senate Majority Leader
511 Dirksen Senate Office Building
Washington, DC 20515

The Honorable Chuck Schumer
Senate Minority Leader
322 Hart Senate Office Building
Washington, D.C. 20510

September 16, 2026

Dear Leaders Thune and Schumer and Members of the U.S. Senate:

As the U.S. Senate works to craft a framework for AI governance, we write to urge you to include critical safety measures that will address the risks that advanced AI systems pose.

Taken together, this summer's wave of multi-agent AI incidents and multiple whistleblower accounts demand a strong response from the U.S. Senate. It should now be clear that the most advanced AI systems present substantial catastrophic and plausibly existential risks to humanity. While AI companies have made progress on voluntary agreements to provide visibility into frontier development, that transparency will continue to fall short if it is not paired with binding federal standards.

Specifically, we strongly request that any AI regulatory framework empower the government, not AI companies, to set mandatory, binding, and enforceable safety standards for advanced AI. Experience has proven time and again that tech companies cannot be allowed to set their own rules when it comes to the safety of their products. With our national security and the future of American communities at stake, we strongly urge you to ensure that forthcoming AI governance legislation meets this minimum requirement.

Leading AI companies have now broken the public trust time and again by releasing products that endanger users. AI companies have not been fully transparent about loss-of-control incidents with unreleased models. Neither the public nor lawmakers have any meaningful insight into what these models can do and whether or not companies are following best practices or even their own commitments to security and safety. While America's AI labs are global leaders in innovation, they cannot be trusted with setting the rules that protect public safety.

In addition to enabling mandatory, enforceable government-set standards for AI, a framework for AI governance should require substantive, transparent, and independent evaluations of advanced AI models and AI development. Recent incidents of AI models escaping human oversight and even committing cybersecurity crimes highlight that we need more examination of models before they are deployed, both internally and externally. Recent voluntary commitments to transparency by leading AI labs are a step forward, but they are ultimately only voluntary and reach only as far as tech companies will allow them to. Auditing of AI models must be both mandatory and conducted independently.

Finally, we strongly urge that Congress establish a strong floor with regard to AI regulations while preserving the critical rights of the states to establish additional regulations beyond those set by Congress. At a moment of rising public concern over quickly advancing AI risks, Congress should not eliminate strong state laws in exchange for weaker federal standards.

The safeguards outlined in this letter are the minimum threshold that we would expect a framework for federal AI governance to meet. A range of other important measures, from whistleblower protections to additional child chatbot safety requirements, deserve Congress's attention. But these standards describe a necessary foundation for a federal AI framework on frontier AI risks.

Thank you for your consideration and for your work to address the issue of AI safety.

Sincerely,

Organizations

The Alliance for Secure AI
Americans for Responsible Innovation
Center for AI Safety Action Fund
Common Cause
CommonDefense.us
Common Sense Media
Consumer Reports
Demand Progress Action
Design It For Us
Economic Security Project Action
Encode AI
FAR.AI
Future of Life Institute
Humans in Control
Institute for Family Studies
Issue One
Public Citizen
Public Knowledge
Secure AI Project
The Tech Oversight Project

Individual Signatories

Alex Bores, New York State Assemblyman, District 73
Malo Bourgon, Machine Intelligence Research Institute
Daniel Kokotajlo, AI Futures Project
Girish Sastry, Guidelight
Alexander Turner, AI researcher