

House Permanent Select Committee on Intelligence
25 Years After 9/11: Confronting the Next Generation of Threats

Testimony of Ben Buchanan

July 21, 2026

Chairman Crawford, Ranking Member Himes, Review Co-Chairs Stefanik and Gottheimer, and Members of the Committee,

Thank you for inviting me here today. My name is Ben Buchanan. During the Biden Administration, I served on the National Security Council and as the first ever White House Special Advisor for AI. I am now a professor at Johns Hopkins University, the author of four books on cybersecurity and AI, and an outside advisor to AI and cybersecurity companies, such as Anthropic.

The 9/11 Commission famously concluded that the attacks revealed a failure of imagination. The government's inability to take seriously a threat that did not fit existing categories and to connect information scattered across institutional seams led to devastating strategic surprise. Twenty-five years later, AI presents, in its own way, a similar kind of challenge. The technology is advancing faster than our institutions are adapting, and the most important data points come from outside government, in private companies. My testimony today is about how our nation and intelligence community can avoid being surprised again.

AI presents substantial risks and benefits to the intelligence community. While in the government, among other things, I worked on National Security Memorandum 25, which aimed to catalyze the use of AI for national security purposes and to prevent strategic surprise to the United States. We had a strong conviction that AI would before long have deep impacts on intelligence collection and analysis, especially in cyber operations.

Since I left the government, it is apparent that this potential is quickly becoming real, and it is a trend that will almost certainly continue. As such, the appropriate governance and use of AI deserve our nation's urgent attention, especially given the rapid pace of technological progress. Let me make three points to ground our discussion.

First, AI is unlike many of the other technologies that are relevant to national security, such as satellites, nuclear weapons, and hypersonic missiles. Those technologies typically originated with government national security programs and developed over decades largely through government direction and oversight. While AI certainly has its roots, decades ago, in government funding, the modern paradigm of large-scale neural networks backed by giant computing power comes almost exclusively from the private sector. This renders government far

less able to understand and shape the technology and increases the risk of strategic surprise. We must together forge a new balanced relationship between the public and private sectors in AI.

Second, we should recognize how powerful AI is *today* and how its capabilities already bear significantly on intelligence operations. Some of the questions that were fiercely debated just a few years ago have now clearly resolved; some of the ideas that were back then hypotheses are now conclusions. For example, AI's capacity to reshape cyber operations was a major focus in the Biden administration and for me personally, and evidence now abounds that it has done so. The nation that can best invent AI and deploy it in offensive and defensive cyber operations will enjoy a distinct advantage over nations that cannot. In my view, it is an imperative that the United States lead the world, as it currently does. We must continue to take actions, such as denying advanced chips and chipmaking equipment to China, that ensure this lead endures.

Again, the threat here is not hypothetical. Both non-state and state actors have tried to use AI's cyber capabilities against the United States and its partners. For example, in November 2025 Anthropic reported that a Chinese state-sponsored group had used its Claude model to run a largely automated cyber-espionage campaign against roughly thirty companies and government agencies across several countries, with the AI carrying out an estimated 80-90 percent of the tactical operation on its own.¹ Google has likewise reported that Russian military intelligence deployed malware in Ukraine that uses AI mid-operation to write its own attack commands.²

And the barrier to entry is falling for criminals, too. In 2025, Anthropic exposed a cybercriminal operation that used AI to hack into and extort a wide range of targets, with the AI automating reconnaissance, credential harvesting, and network penetration. Other criminals have used AI to develop and distribute ransomware software.³

Third, we should expect AI to rapidly get much better. The most important input to AI capability is computing power, and the largest U.S. technology companies plan roughly \$700 billion in capital spending this year, overwhelmingly on AI infrastructure.⁴ This amount is more than the United States spent on the Interstate Highway System in 35 years (adjusting for

¹ Anthropic, "Disrupting the First Reported AI-Orchestrated Cyber Espionage Campaign," November 13, 2025, <https://www.anthropic.com/news/disrupting-AI-espionage>.

² Google Threat Intelligence Group, "GTIG AI Threat Tracker: Advances in Threat Actor Usage of AI Tools," November 5, 2025, <https://cloud.google.com/blog/topics/threat-intelligence/threat-actor-usage-of-ai-tools>. On the attribution to the GRU, see Cybersecurity and Infrastructure Security Agency, National Security Agency, Federal Bureau of Investigation, et al., "Russian GRU Targeting Western Logistics Entities and Technology Companies," Joint Cybersecurity Advisory AA25-141A, May 21, 2025, <https://www.cisa.gov/news-events/cybersecurity-advisories/aa25-141a>.

³ Anthropic, "Threat Intelligence Report: August 2025," August 27, 2025, <https://www.anthropic.com/news/detecting-counteracting-misuse-aug-2025>.

⁴ Vlad Savov, "US Big Tech Ratchets Up AI Spending Past \$700 Billion This Year," *Bloomberg*, April 30, 2026, <https://www.bloomberg.com/news/articles/2026-04-30/us-big-tech-ratchets-up-ai-spending-past-700-billion-this-year>.

inflation) and akin to a Manhattan Project every two or three weeks. More powerful systems will follow, especially as companies develop more efficient algorithms.

In addition, AI systems are increasingly helping AI developers train their successors – a major step known as recursive self-improvement, and one of the most important ideas in modern AI. While full self-improvement remains ahead of us, this process has clearly begun. For example, in early 2026 OpenAI released a model it called GPT-5.3-Codex, describing it as “our first model that was instrumental in creating itself.” The company said it had used earlier versions to debug its own training, manage its own deployment, and diagnose its own test results.⁵ Anthropic has similarly said that the vast majority of its production code is written by AI but overseen by humans. As a result, the company says that its typical engineer produces eight times as much code in 2026 as in 2024.⁶ Independent analyses by the nonprofit METR find that the length of tasks that leading AI systems can complete on their own has been doubling every few months; the best models can now handle tasks that take a skilled person several hours.⁷

Putting these three things together – government unfamiliarity, proven national security relevance, and strong reasons to expect rapid future progress – should yield a sense of stark urgency. AI will enable a lot of good, but it will also raise new and grave concerns in areas as diverse as cybersecurity, biology, and accidental loss of control. The Committee is right to focus its attention on this technology, and I am honored to help you however I can in this endeavor.

⁵ OpenAI, “Introducing GPT-5.3-Codex,” February 5, 2026, <https://openai.com/index/introducing-gpt-5-3-codex/>.

⁶ The Anthropic Institute, “When AI Builds Itself,” June 4, 2026, <https://www.anthropic.com/institute/recursive-self-improvement>.

⁷ METR, “Time Horizon 1.1,” January 29, 2026, <https://metr.org/blog/2026-1-29-time-horizon-1-1/>.