

AI is a cybersecurity problem — and also, a solution

ASU expert Nadya Bliss shares takeaways from her new report on the future of AI and cybersecurity

By Mikala Kass, ASU News
June 26, 2026

Artificial intelligence is changing cybersecurity. It's a two-edged sword that can help both cyber criminals and cyber defenders. While the bad guys seem to hold the advantage for now, we can turn the balance in the good guys' favor by investing in better, AI-supported cybersecurity practices.

That's the opinion from a new "[rapid expert consultation](#)" report co-authored by [Nadya Bliss](#), executive director of the [Advanced Capabilities for National Security Institute](#) at Arizona State University. The rapid expert consultation is a product of the National Academies of Sciences, Engineering, and Medicine.

"The way that I tend to think about it is, it used to be that you had to be a pretty sophisticated attacker to launch a sophisticated attack, and now that is no longer true," Bliss says. "This is not something we can sweep under the rug. We have to address this shift to protect our digital systems."

Below, Bliss answers questions about what AI means for our daily lives as well as our national security landscape.

Note: Answers edited for length and clarity.

Question: What are the key takeaways from the report that came out this week?

Answer: A big takeaway from the report is that AI is fundamentally transforming cybersecurity. In the short term, AI is likely to benefit the attacker, just because of the nature of the beast. An attacker only has to be right once, and a defender has to be correct all the time. But in the long term, we are incredibly hopeful that this is an opportunity to have more secure systems and give more users tools to protect their systems automatically.

Q: The average household has bank accounts, passports and private medical information moving around in these software environments. What is your best advice for this current

moment?

A: If you think about how attackers operate, they essentially look for vulnerabilities in the systems. Sometimes those vulnerabilities are machines; sometimes those vulnerabilities are humans. Both of those modalities of attack are now significantly enabled by artificial intelligence. I do think high-profile organizations like banks are aware of some of these vulnerabilities. I have noticed recently that they have strengthened their defenses and emphasized things like two-factor authentication and passkeys. All of the advice for individuals that we have given in the past about making sure not to click on things, making sure not to give passwords or share information over the phone, that still very much applies.

Q: What is the margin of time between this moment of concerning vulnerability and the moment when our countermeasures catch up?

A: This is an important aspect to consider. In the report, we argue that in the near term the attacker is advantaged, and in the longer term we think that the defender will be advantaged. Trying to compress the length of time between those two states is precisely what we're advocating for. How well we do that depends on whether we have effective coordination, effective public-private partnership and an appropriate set of incentive structures, including investments to build out those defenses. We need defenders to leverage AI across their systems, just as attackers could now do pretty readily.

Q: Are we looking at an evolution where AI is the problem, and yet AI itself could ultimately be the solution?

A: The way that I tend to think of technology — and AI is just a type of technology — is that it can be used for good or bad. There is an important parallel between this moment in time with artificial intelligence and what we experienced as a society in the '90s and early 2000s, a period when capability developed way faster than any guardrails around that capability. The capability by itself is not inherently bad or good; it's just the capability. But we need to build out guardrails in an efficient way to make sure that we benefit from those capabilities as opposed to become victim to attackers misusing them.

Q: Is there any connection between the broad strokes of your findings and the recent big headlines we've heard about the capabilities of Anthropic's Claude Mythos model?

A: The frontier AI model companies are developing and deploying capabilities at an incredibly fast rate, and there's a number of those companies. Mythos' capability was developed and discussed right while we were developing our rapid expert consultation. So that's a good example to look at — both in terms of mitigating risks and assessing capabilities.

I will say that Mythos was initially limited in where it got released precisely because the type of capability it provided could be dangerous. We as authors feel that it is not sufficient to just limit the release of technology. It's much more important to build out systemic resilience and what we call "defense-in-depth" longer term. In other words, we need to develop a robust, adaptable, persistent cybersecurity ecosystem.

Q: What is the scariest development you're seeing, and what brings the biggest sense of relief around that anxiety?

A: I was at the very beginning of my computer science career in the late 1990s and early 2000s. At the time, to me it was obvious that we were creating and deploying systems that were vulnerable. That is when the internet became a household thing and everybody started participating in social media. I remember thinking, “There are so many holes in all of this.” Data breaches were an obvious risk; negative impacts from social media seemed like an obvious risk. It took a number of pretty significant negative consequences for some of those vulnerabilities to be curtailed.

We're a lot more secure now. There's more infrastructure on social media to protect users. There's more infrastructure on interconnected systems to protect users. Some of that infrastructure is technological, some of it is policy, some of it is incentive-based. What I am hoping is that we have learned from those mistakes, and we're not going to repeat them with artificial intelligence. Is there tremendous capability and tremendous hope and optimism? Yes, but it has to be done with open eyes, understanding what the risks are. We tend as a society to overfocus on the capability and underfocus on security. Things are moving a lot faster, but we also know a lot more. So let's do this better than we did with the internet.

Q: Are we in a world now where we need to have assessments like this rapid expert consultation continuously?

A: Absolutely. I think there is a place for both longer-term assessments and rapid assessments. The reason that I would recommend a continuous rapid assessment of technology, at least in this particular moment in time, is because the diffusion of artificial intelligence is at an unparalleled scale. What's interesting, especially about the generative model type of artificial intelligence, is that even experts who study it often cannot tell you precisely why it works.

If the experts can't explain how things work and we are giving it to every single user out there, that creates a significant gap between understanding and usability.

Having a continuous reassessment of the implications of AI on various industries — AI and science, AI and health care, AI and banking, AI and travel, AI and entertainment, AI and the creative arts — and more broadly the impact of AI on society as well, I think those are other areas where we need expert consultations.

Q: Does AI have implications for national security and defense?

A: AI is central to national security and defense. That's not just me talking — the Pentagon has been aggressively pursuing AI adoption measures and AI capability advancements. This is both to enable the warfighter and to protect them from adversaries using AI for their own advantage. The implications for national security are countless: from protecting critical infrastructure in energy, health care and our water supply, to maintaining our ability to operate in contested environments. AI is necessary for all these capabilities.

This is an area of strength at ASU — applying AI advancements for national security enhancement. We have active projects working on using AI to bolster hospital cybersecurity, to improve military training performance and to increase communications speed between space-based assets.

Steve Filmer and Mikala Kass contributed to this reporting.

This story originally appeared on [ASU News](#).

Main image



Nadya Bliss is the executive director of the Advanced Capabilities for National Security Institute at ASU, as well as a professor of practice in the School of Computing and Augmented Intelligence. Photo illustration by Andy DeLisle and Christian Van Bebber/ASU