

The Skill-Ability Gap: Why Five Eyes' AI Warning Is Systemic

2026-07-12

 AI-assisted Rapid Research



CSAI

cloud
CSA security
alliance®

© 2026 Cloud Security Alliance. Some rights reserved.

You may download, store, display, view, print, redistribute, and link to this document in its original, unmodified form, provided that attribution to the Cloud Security Alliance is maintained and all trademark and copyright notices remain intact.

This document may not be modified or altered. You may quote portions of the document as permitted by the Fair Use provisions of the United States Copyright Act, provided that attribution is given to the Cloud Security Alliance.

This document may be shared on professional and social media platforms in its original form with attribution.

This document was generated with AI assistance and has not undergone official CSA review and approval processes.

Key Takeaways

On June 23-24, 2026, the cyber security agencies of the Five Eyes nations – the United States, United Kingdom, Canada, Australia, and New Zealand – issued a rare joint statement warning that frontier AI models will "fundamentally transform both offensive and defensive cyber capabilities," on a timeline the agencies described bluntly: "The timeline is not years, it is months" [1][2]. The statement followed weeks after the U.S. government ordered Anthropic to suspend foreign-national access to its Mythos 5 and Fable 5 models over concerns about their vulnerability-discovery capabilities – an order Euronews reported was directed by President Trump [2][3]. Security researcher Bruce Schneier situated the warning inside a longer argument he had been building since June: AI is accelerating a decades-old process in which "skill" and "ability" – once synonymous in cybersecurity – have become decoupled, so that people with no specialized training can now execute attacks that once required years of expertise [5][6]. This matters because the same 2025 ISC2 Cybersecurity Workforce Study that underlies most industry planning had already abandoned headcount-gap estimates in favor of skills-gap measurement, finding that 59 percent of practitioners report critical or significant skills gaps and 88 percent experienced a security event attributable to a skills shortage in the prior year [7]. In CSA's reading, these developments together indicate that the Five Eyes advisory is not a warning about a discrete new threat actor or exploit, but a recognition that a chronic, already-documented structural weakness in cyber defense – insufficient skill relative to the ability organizations need to defend themselves – is being exploited at a pace that outstrips how quickly institutions can respond.

Background

The Five Eyes statement was signed by the heads of each nation's principal cyber security authority: Stephanie Crowe, head of the Australian Cyber Security Centre; Rajiv Gupta, head of the Canadian Centre for Cyber Security; Catriona Robinson, head of New Zealand's National Cyber Security Centre; Richard Horne, CEO of the UK's National Cyber Security Centre; David Imbordino, director of the U.S. National Security Agency's Cybersecurity Directorate; and Nick Andersen, acting director of the U.S. Cybersecurity and Infrastructure Security Agency [1]. The advisory's central claim is that malicious actors are already "using AI to shorten the time between the discovery of a vulnerability and its exploitation," a dynamic the agencies argue will only accelerate as frontier models improve [1]. Their recommended response leans on fundamentals rather than novel controls: understand and assess risk, prioritize

foundational cyber hygiene, patch and strengthen access controls faster, and – notably – "give cyber leaders the authority and resources needed to respond as the threat landscape changes" [1][2]. The agencies were explicit that "those who delay will face growing and avoidable risk" [1]. Former CISA director Chris Krebs, reacting to the statement, called it "pretty alarming" and warned of an approaching "vulnerability tsunami" as AI-assisted discovery outpaces organizations' capacity to remediate what it finds [4].

The advisory did not arrive in a vacuum. In April 2026, Anthropic disclosed that its Mythos-class models demonstrated what CBS News characterized as "unprecedented abilities to find software vulnerabilities" [4]. That disclosure prompted the U.S. government to order Anthropic, in June 2026, to suspend Fable 5 and Mythos 5 access for foreign nationals [2][3]. Canada, which had gained early access to Mythos 5 to test government systems, lost that access under the same directive. Prime Minister Mark Carney responded by framing the episode as a lesson in overreliance, saying "the situation we're in collectively right now with Mythos and Fable is something that can happen with overreliance on certain models," and pointing toward the need for Canada to build out and diversify its own AI capacity [3].

It was against this backdrop that Schneier published two connected essays reframing the moment. In "Once, Cyberattacks Required Great Skill. AI Is Changing That," he traced a lineage from the seven L0pht hackers who testified to Congress in 1998 that they could disable large parts of the internet within thirty minutes, through the rise of "script kiddies" who ran attack tools built by more capable others, to today's AI systems that can "with little detailed direction, autonomously hack into networks, steal data, deploy ransomware" [6]. His follow-up essay, "Cybersecurity and the Gap Between Skill and Ability," generalized the argument: the long apprenticeship once required to gain dangerous capability also functioned, almost incidentally, as an ethical filter, since acquiring deep technical skill typically meant absorbing the norms of a professional community along the way [5]. AI severs that link. A model can grant someone the ability to exploit a vulnerability without requiring them to acquire the skill, and therefore without the socialization that historically accompanied it – a dynamic Schneier extends beyond cybersecurity to note that doctors, virologists, and structural engineers hold comparably dangerous knowledge that has likewise depended on professional gatekeeping rather than technical barriers.

Security Analysis

In CSA's assessment, this dynamic is difficult to resolve through model-level refusal alone, because the same capability that enables offense is the capability the Five Eyes advisory itself recommends organizations adopt for defense. The agencies explicitly urged organizations to use AI to "detect vulnerabilities earlier, improve software quality, monitor unusual behaviour, and respond faster to incidents" – the identical model behavior, applied to identical code, that enables faster exploitation in

the hands of an attacker [4][6]. Vendor-side guardrails and export controls appear, in CSA's assessment, to address this only at the margins. Frontier labs can restrict their own hosted models, as Anthropic did under government order, but Schneier's point about smaller open-weight models proliferating "like script kiddie hacker tools" without comparable safeguards means the underlying capability continues to diffuse regardless of what any single vendor restricts [5][6]. The Anthropic episode is therefore better understood as a demonstration of the limits of vendor-level containment than as a solution: it removed one access path for foreign nationals to two specific models while leaving the broader capability trend largely unaddressed [3].

This is where the workforce data becomes load-bearing rather than background context. ISC2's 2025 Cybersecurity Workforce Study, drawing on more than 16,000 practitioners globally, marked a deliberate shift away from the workforce-gap framing the industry had used for years, because respondents increasingly identified insufficient skill, not insufficient headcount, as the binding constraint [7]. Ninety-five percent of respondents reported at least one skill need, up five points year over year, and 59 percent cited critical or significant skills gaps, up from 44 percent – a 15-percentage-point increase, or roughly 34 percent in relative terms – even as 88 percent said their organization had experienced a security event in the past year attributable to a skills shortage [7]. That data was gathered before the Five Eyes advisory, but it describes precisely the asymmetry the advisory now warns is being exploited: defenders already lacked sufficient specialized skill relative to what their environments demanded, and AI is not creating that gap so much as handing attackers a shortcut around the equivalent gap on their side, while leaving defenders' gap fully intact. An attacker who previously needed years of study to weaponize a disclosed vulnerability can now lean on AI assistance to close much of that gap; a defender facing that same vulnerability still needs the skill to triage, patch, and validate the fix, and that skill has not gotten any easier to acquire.

The compression of exploitation timelines echoed in Krebs's "vulnerability tsunami" framing risks turning pre-existing structural weaknesses – technical debt, deferred patching, unmanaged legacy systems – from chronic liabilities into acute ones [4]. Organizations that have tolerated slow patch cycles because exploitation windows were historically measured in weeks now face the same unpatched systems under exploitation windows that AI-assisted discovery is compressing toward days or hours. In CSA's reading, the agencies' recommended remedies are established cyber hygiene practices rather than novel controls, which supports treating this as a structural rather than a discrete threat [1][2]. That is itself evidence for the systemic reading of this event. If the fix required were a new technology or a new control, this would be an incident calling for an incident response. Because the fix required is the same foundational discipline security teams have been underfunded and understaffed to execute for years, this is instead a long-standing structural weakness that AI has made acutely visible rather than newly created.

The table below summarizes how the barrier to executing a serious cyberattack has shifted across three eras, illustrating why the current moment represents a step change rather than a continuation of the script-kiddie dynamic security teams have managed for two decades.

Era	Representative Actor	Barrier to Entry	Ethical/Professional Filter
Skilled hacker (pre-2000)	1998 L0pht witnesses testifying to Congress [6]	Years of self-taught or professional technical training	Strong – expertise typically acquired alongside community norms
Script kiddie (2000s-2010s)	Users of pre-built exploit tools and kits	Access to tools built by more skilled others; minimal technical depth required	Weak – tool use required no socialization into a professional community
AI-enabled (2026-)	Any user of a capable frontier or open-weight model	Natural-language prompting; little to no technical background required	Effectively absent for models without vendor guardrails, and appears to be diminishing even where guardrails exist [3][5][6]

Recommendations

Immediate Actions

Security leaders should treat the Five Eyes advisory's compressed exploitation timeline as an input to patch prioritization rather than as general commentary, accelerating remediation schedules for internet-facing systems and reviewing whether existing service-level targets for patching still make sense when exploit development can plausibly occur far faster than in the past. Organizations should also inventory legacy and end-of-life systems that have persisted under an assumption of slow exploitation timelines, since these are the assets whose risk profile changes most sharply under AI-compressed discovery-to-

exploit windows. Consistent with the advisory's explicit recommendation, executives should confirm that cyber leaders hold the authority and budget to act quickly on emerging risk without escalation delays, since the agencies specifically flagged authority gaps as a factor that slows response [1].

Short-Term Mitigations

Enterprises should begin adopting AI-assisted defensive tooling for vulnerability detection, code review, and anomaly monitoring, matching the same class of capability that the advisory warns is being used offensively, rather than treating AI adoption in security operations as optional or aspirational [1][2]. At the same time, organizations should avoid over-relying on any single AI vendor's guardrails or export-control compliance as a durable safety control, since capability restricted at one vendor can persist across the broader open-weight and third-party model ecosystem [5][6]. Because the underlying workforce problem is a skills gap rather than a headcount gap, security teams should redirect a portion of hiring budget toward structured upskilling – particularly in AI-assisted triage, secure code review, and incident response under compressed timelines – rather than assuming additional headcount alone will close the exposure ISC2's data describes [7].

Strategic Considerations

Boards and executive leadership should adopt the advisory's framing that cybersecurity is a core business risk and leadership responsibility rather than a purely technical function, and should expect that framing to persist well beyond this specific advisory cycle [1][3]. Organizations with meaningful dependence on a small number of frontier AI vendors should treat the Anthropic export-control episode as a precedent rather than an anomaly, and should build contingency plans for sudden, government-directed loss of access to AI capabilities that have become embedded in security operations [3][4]. Finally, because the skill-ability gap Schneier describes is structural rather than incident-specific, organizations should plan multi-year investment in workforce skill development and AI-literacy training for security staff, recognizing that no single advisory response, patch cycle, or vendor restriction will resolve a gap rooted in how broadly capability has diffused relative to how narrowly skill remains distributed [5][6][7].

CSA Resource Alignment

CSA's own [“The AI Vulnerability Storm”: Building a “Mythos-ready” Security Program](#) [8], published in April 2026, engages directly with the dynamic this note's Background section describes: AI models compressing the interval between vulnerability discovery and exploitation, and the operational

adjustments a security program needs to keep pace with that compression. Organizations acting on the Five Eyes advisory's recommendation to tighten patch cycles and adopt AI-assisted detection should treat this report as the more narrowly tailored starting point, given its direct focus on the same compressed-timeline dynamic this note traces to a policy level.

CSA's [Using AI for Offensive Security](#) [9] research report, published in 2024, is also directly relevant: it examines how large language models and AI agents compress reconnaissance, vulnerability analysis, and exploitation across the offensive security lifecycle. In CSA's reading, the report's account of AI lowering the skill threshold for offensive testing anticipates the same phenomenon Schneier later described as the decoupling of skill from ability, and organizations building internal governance around AI-assisted offensive tooling should treat it as a starting framework for scoping that governance.

Because the Five Eyes advisory centers on AI systems that can act with minimal human direction to autonomously identify and exploit vulnerabilities, CSA's [MAESTRO agentic AI threat modeling framework](#) [10] provides a structured way to reason about the autonomy dimension of this risk specifically, beyond the offensive-technique focus of the reports above. Security teams evaluating how much autonomous latitude to grant AI-assisted defensive tooling – the same category of tool the advisory recommends adopting – can use MAESTRO's threat categories to ensure that defensive automation does not introduce its own unmanaged agentic risk while organizations race to keep pace with AI-accelerated offense.

More broadly, the governance questions raised by this episode – how organizations should manage dependency on frontier AI vendors, verify safeguard claims, and integrate AI capability into security operations responsibly – map onto the control domains in CSA's [AI Controls Matrix \(AICM\) v1.1](#) [11]. Organizations responding to the Five Eyes advisory by adopting AI-assisted security tooling, as the agencies recommend, should use AICM's control domains as the structural backbone for the governance program that adoption requires, rather than building bespoke oversight from scratch under time pressure.

References

- [1] SecurityBrief Australia. "[Five Eyes warn AI cyber risks are rising within months](#)." SecurityBrief Australia, June 24, 2026.
- [2] Euronews. "[AI cyber threat is 'months, not years' away, Western intelligence agencies warn](#)." Euronews, June 23, 2026.
- [3] Global News. "[Five Eyes issue 'call to action' as AI becomes a 'core' cybersecurity risk](#)." Global News, June 24, 2026.
- [4] CBS News. "[AI on pace to bypass cybersecurity systems in months, not years, 'Five Eyes' spy partners warn](#)." CBS News, June 23, 2026.
- [5] Schneier, Bruce. "[Cybersecurity and the Gap Between Skill and Ability](#)." Schneier on Security, July 2026.
- [6] Schneier, Bruce. "[Once, Cyberattacks Required Great Skill. AI Is Changing That](#)." Schneier on Security, June 2026.
- [7] ISC2. "[2025 ISC2 Cybersecurity Workforce Study](#)." ISC2, December 2025.
- [8] Cloud Security Alliance. "['The AI Vulnerability Storm': Building a 'Mythos-ready' Security Program](#)." CSA, April 2026.
- [9] Cloud Security Alliance. "[Using AI for Offensive Security](#)." CSA, 2024.
- [10] Cloud Security Alliance. "[Agentic AI Threat Modeling Framework: MAESTRO](#)." CSA, February 2025.
- [11] Cloud Security Alliance. "[AI Controls Matrix \(AICM\) v1.1](#)." CSA, 2026.