



Zscaler Joins the Hiroshima AI Process Partners Community

July 20, 2026

Zscaler has joined the [Hiroshima AI Process \(HAIP\)](#) Friends Group Partners' Community, reinforcing our commitment to advancing secure, transparent, and trustworthy AI adoption.

Launched under Japan's 2023 G7 Presidency, the HAIP advances international discussion on safe, secure, and trustworthy AI. Its Partners Community brings together organizations committed to supporting practical progress on AI governance and responsible adoption.

As AI adoption accelerates, organizations are not only deploying AI more broadly but also confronting a growing range of AI-powered attacks. Zscaler AI helps them use AI safely for productivity, protect their AI initiatives, and improve their security outcomes with AI-driven defenses and controls.

Zscaler will contribute to the HAIP our lessons from this real-world experience securing AI and defending against AI-enabled attacks. Importantly, the HAIP already provides a framework for this through its International Guiding Principles and Code of Conduct for advanced AI systems, which emphasize risk assessment, security testing, transparency, incident response, and safeguards against misuse.

This focus aligns with Zscaler's approach to AI security. Alongside support for international governance efforts, Zscaler is actively engaged in frontier AI security initiatives including Anthropic's Project Glasswing and OpenAI's Trusted Access for Cyber program. Participation in these efforts strengthens understanding of advanced AI capabilities in cybersecurity and supports more rigorous security testing and evaluation.

Why transparency matters

Trust in AI cannot rest on broad assurances alone. It depends on evidence, visibility, and candid reporting.

That principle is reflected in Zscaler's approach to AI security research and threat intelligence publications. Our ThreatLabz team regularly publishes AI security findings, including an annual AI Security Report. [The 2026 report](#) found that AI adoption is creating critical security gaps across global enterprises.

Zscaler has also moved quickly to share its frontier model security research. In "[When the Scanner Starts Thinking: Learnings from Mythos and GPT 5.5 Cyber Security](#)," Zscaler shared a structured evaluation of frontier models in cybersecurity, outlining our testing methodology, assessment of model capabilities, and practical security recommendations.

Zscaler has also been transparent about the company's internal AI governance, including its commitment not to use proprietary customer data or personal information to train AI models. Sam Curry, Zscaler's Global CISO, has [explained this commitment](#) and our data containment architecture and use of anonymized, aggregated signals to improve security while protecting privacy. These actions reflect Zscaler's commitment to using AI in a way that is governed, transparent, and responsible.

A global commitment

Joining the Hiroshima AI Process Partners Community is an important milestone in Zscaler's broader commitment to supporting secure digital transformation worldwide.

As AI policy and practice continue to evolve, Zscaler looks forward to contributing its perspective on how security, transparency, and practical governance can help organizations securely adopt AI and defend against its malicious uses.

Zscaler also looks forward to working with governments, industry partners, and the broader Hiroshima AI Process community to help ensure that AI adoption is matched by practical safeguards, stronger transparency, and security by design.