

Netskope Threat Labs Report

Australia & New Zealand

September 2026

The 2026 Netskope Threat Labs Australia & New Zealand (ANZ) report details the increasing adoption of AI, trends in data policy violations, malware distribution via cloud applications and phishing trend observed over the last year.

Key findings

This report examines how AI is reshaping the digital footprint of organizations across ANZ and the security and data-protection challenges that come with it. It also discusses how attackers are abusing trusted cloud platforms to push malware into local organizations, and explores how the current patterns of data policy violations should serve as a call to action for organizations handling sensitive data.

Shadow AI is losing ground: Staff using personal AI accounts and sending potentially sensitive work-related data in prompts, has been a major AI security challenge in recent years. In the last twelve months, organizations in ANZ have made significant ground in curbing this Shadow AI vector, with usage rates of organization-managed AI tools more than doubling over the period, leading to a drop in personal applications' usage. But as employees refine their preferences in terms of AI usage, and continue to explore new AI tools and applications, the amount of users bouncing between personal and enterprise accounts almost doubled as well, highlighting the need for organizations to implement frictionless processes for employees to submit new AI tools and applications for review and approval, and steer them away from creating personal accounts.

Anthropic has pulled ahead of ChatGPT: While ChatGPT remains the most used AI application in the large majority of organizations in the regions around the world that we monitor, Anthropic's Claude is well in front in ANZ, having experienced a massive increase in adoption in just a few months. This trend is also reflected in AI API adoption, where Anthropic's API has established a significant lead over competing providers.

Direct AI usage is just the tip of the iceberg: Nearly all employees now use software with embedded AI features, or features that feed user data into models for training purposes. It permeates workflows and operations, often in ways that do not manifest in direct usage, making discovery, monitoring, and governance all the more challenging.

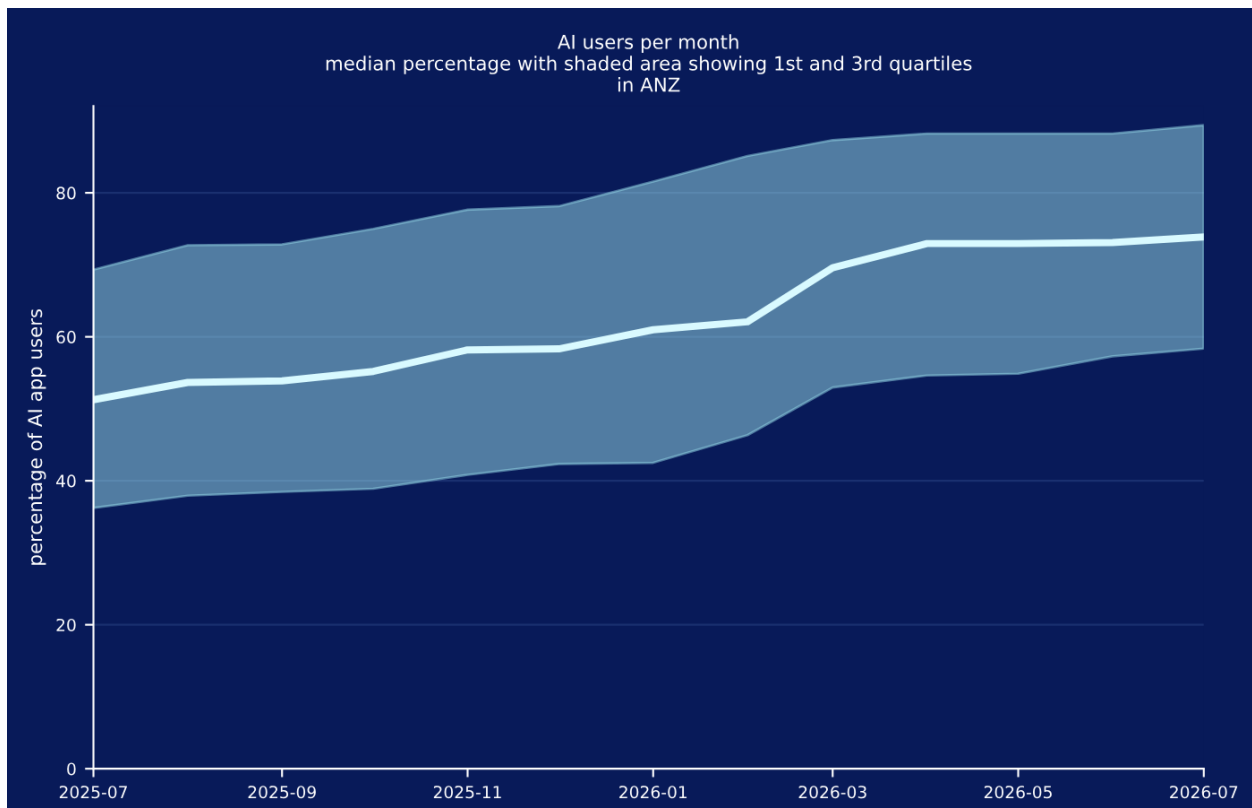
Regulated data is the most at risk: In ANZ, regulated data is the category most exposed in data policy violations, ahead of intellectual property and source code. The type of information supposed to be the most tightly controlled is precisely the one most likely to leak. These findings highlight the need for solid DLP controls and clear AI rules.

Model Context Protocol (MCP) traffic grows quickly as AI and agentic workflows multiply: MCP is the open standard that lets AI models interact with data sources and tools, and MCP traffic is rising rapidly in ANZ. Much of this momentum is driven by the use of coding assistants like Claude Code and Codex. While MCP traffic is a sign that agentic AI adoption is progressing fast within organizations across ANZ, each of these connections is a new pathway for company data to move between AI apps and outside systems, and a vector of potential data leaks if not properly governed.

AI use

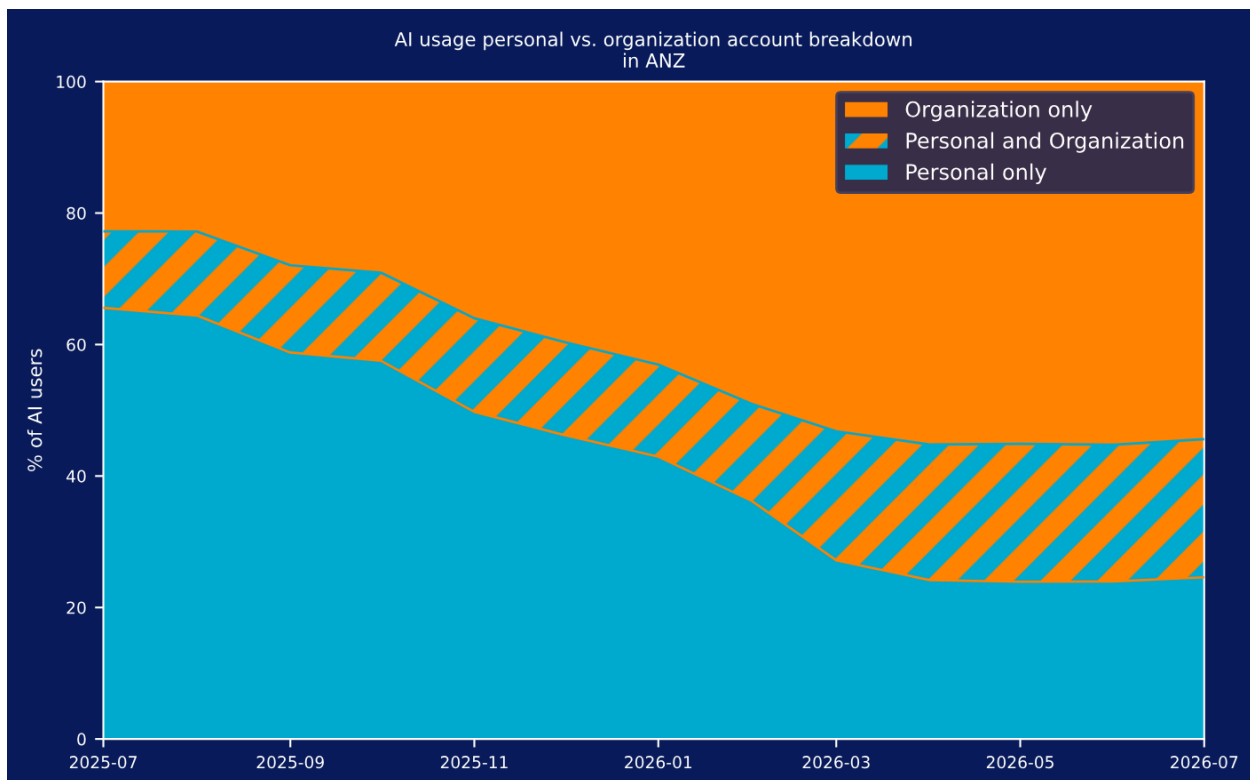
AI: Adoption and usage trends

As organizations grow more confident with adopting and using AI, the technology is becoming a key component of daily business operations. AI use across ANZ has continued to grow over the past year, with the share of users actively using AI applications increasing from **53% to 75%**.

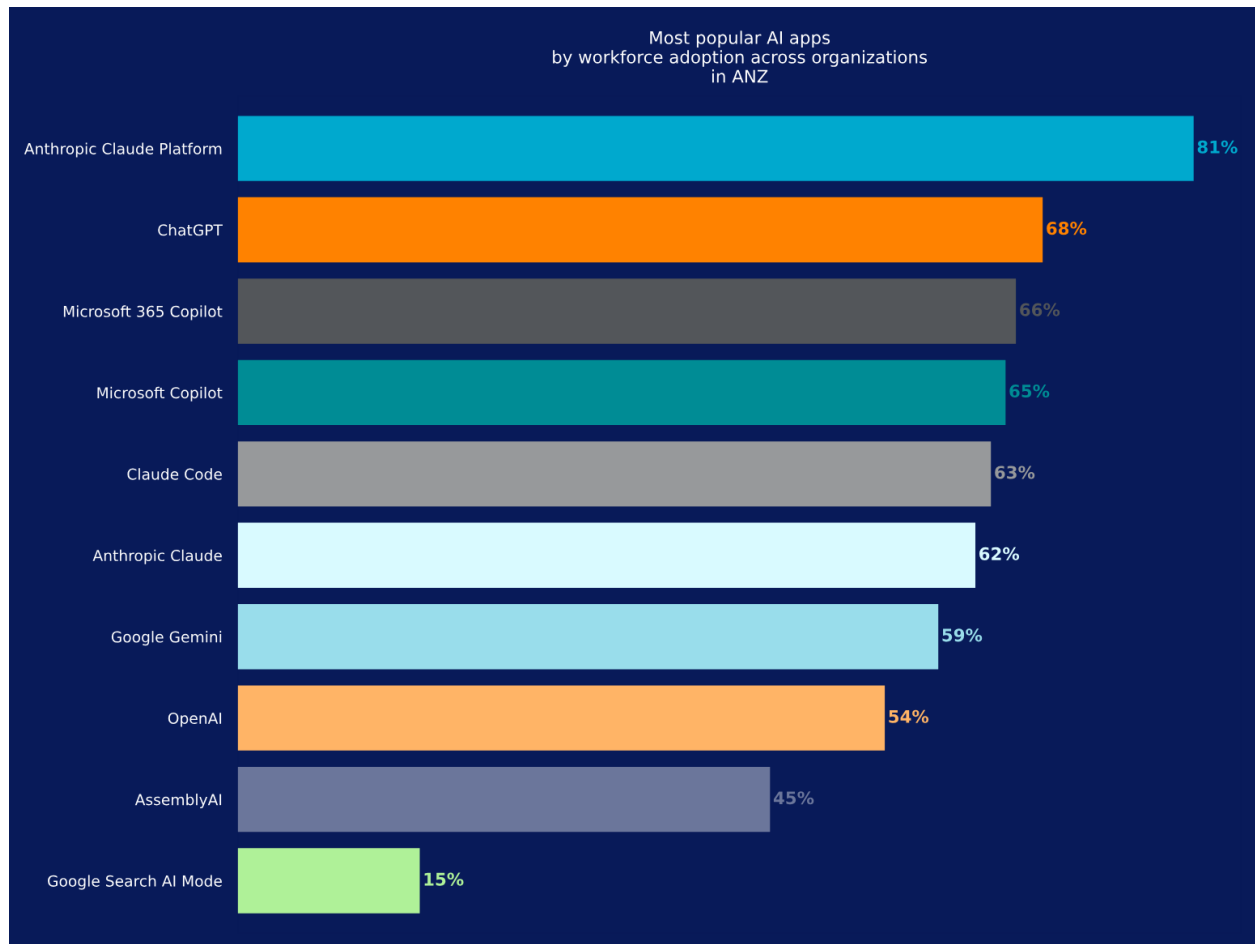


In parallel, organizations across ANZ have made steady progress in reducing shadow AI risk by shifting users away from personal AI accounts and toward organization-managed tools. Over the past year, the use of personal AI applications declined from 76% to 55%, while adoption of organization-managed AI solutions more than doubled, rising from 34% to 75%. At the same time, the share of users switching between personal and enterprise accounts increased from 11% to 21%, suggesting that many employees continue to switch between managed and personal AI accounts as they explore new AI tools, highlighting the need for organizations to streamline the review and approval of new AI applications.

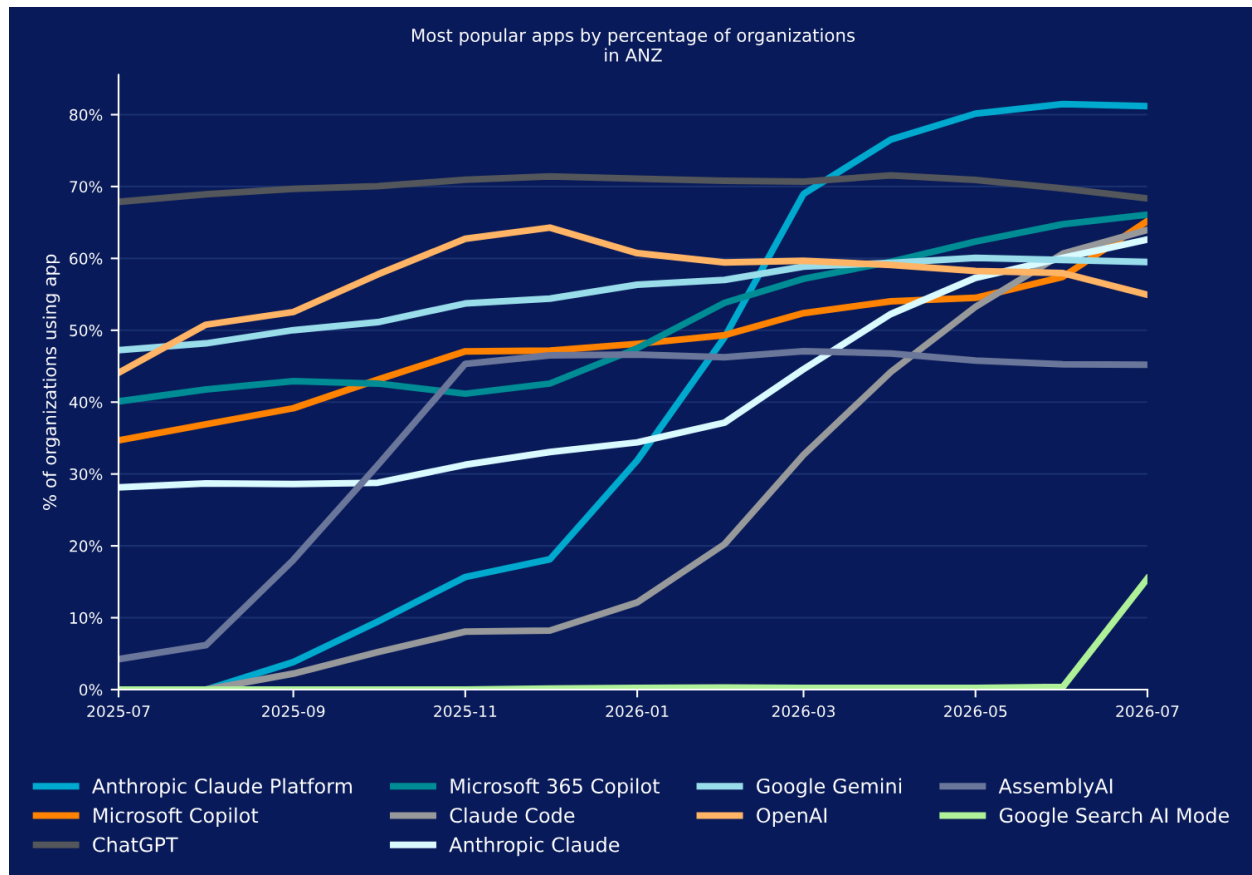
Overall, these changes point to stronger oversight and a broader move toward managed AI deployments that improve data protection and compliance without limiting productivity. Even so, the continued overlap between personal and enterprise usage shows there is still work to do to address shadow AI fully.



AI adoption patterns differ noticeably from global trends. **Anthropic Claude Platform has overtaken ChatGPT as the most widely adopted AI application in ANZ, and is now used by 81% of organizations**, compared to 68% for ChatGPT and 66% for Microsoft 365 Copilot. Unlike the global trend, where ChatGPT continues to lead, organizations in ANZ appear to be placing greater emphasis on AI platforms designed for enterprise use, reflecting evolving preferences as AI adoption matures.



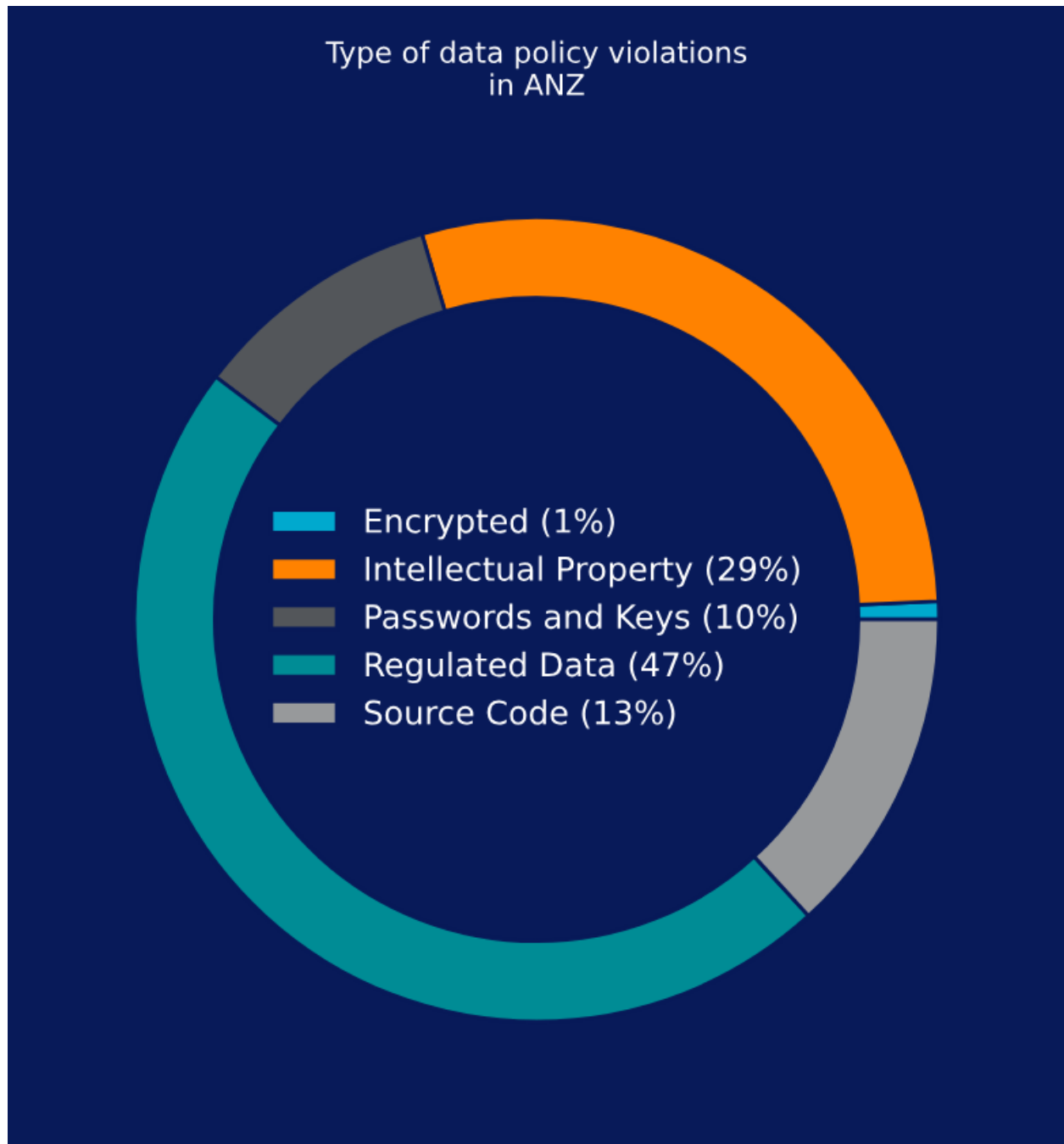
The chart below shows how adoption of the leading AI applications has evolved across ANZ over the past year, revealing a noticeable shift in application preference. While ChatGPT remained one of the most widely used AI applications throughout the period, Anthropic Claude Platform experienced a sharp surge in adoption beginning in December 2025. Its growth accelerated rapidly over the following months, until it overtook ChatGPT around March 2026 as the most used AI application in the region. Meanwhile, Microsoft 365 Copilot has remained relatively stable, with adoption increasing at a more gradual pace.



AI: App usage and data policy violation

As AI adoption continues to grow across organizations in ANZ, so do concerns around potential data exposure. AI tools are now widely used to summarize documents, generate reports, and support day-to-day business processes, many of which involve sensitive business or customer information, expanding the potential attack surface. In this environment, protecting data remains a top priority, particularly as organizations continue to address shadow AI risks.

Analysis of AI-related data policy violations across ANZ shows that regulated data accounts for the largest share of attempts to leak sensitive information in AI prompts at 47%, followed by intellectual property at 29%, source code at 13%, passwords and API keys at 10%, and encrypted data at 1%.

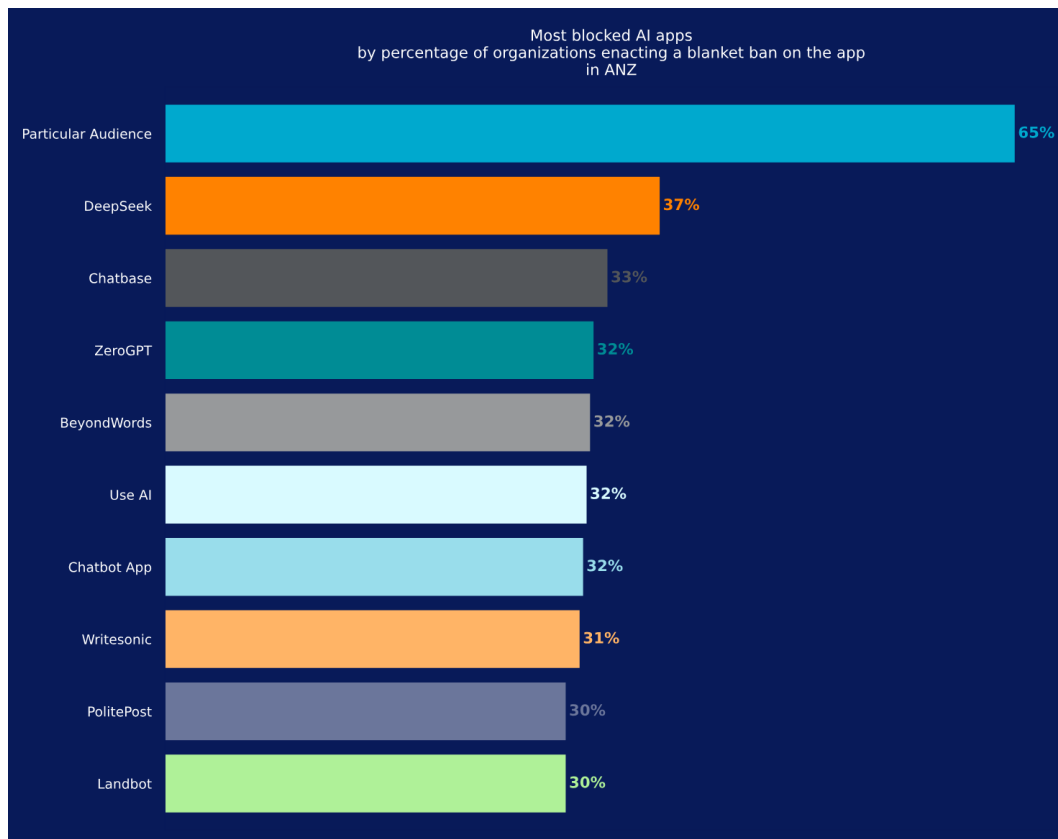


Most blocked AI apps

Organizations across ANZ are taking a cautious approach to AI adoption, often restricting certain applications due to security, privacy, and compliance concerns. While policies vary from one organization

to another, certain applications are blocked more frequently than others, reflecting where the highest risks are perceived.

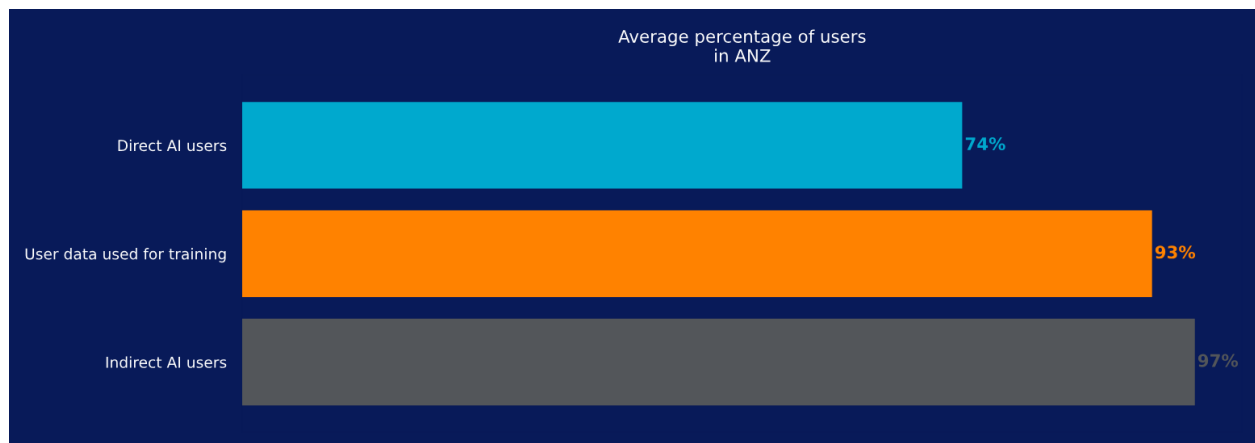
Particular Audience is the most frequently blocked AI application, restricted by 65% of organizations in ANZ. The platform is designed for online retailers, and can involve processing sensitive business or customer information. It is followed by DeepSeek at 37%, Chatbase at 33%, and ZeroGPT at 32%. Chatbase is widely used to build AI-powered chatbots on top of proprietary business data, while ZeroGPT is commonly used to detect AI-generated content. These applications can increase data exposure risk because they interact directly with internal data sources that in many cases can include sensitive information. Overall, these patterns suggest that organizations in ANZ are strengthening governance by proactively limiting the use of AI applications that present higher data security or compliance risks.



User adoption of AI

AI adoption now spans multiple layers: In ANZ, 74% of employees use AI applications directly, while 97% use applications that incorporate AI-powered features. In addition, 93% interact with AI systems that use customer or user data to train models.

These figures illustrate how deeply AI has become embedded in everyday work, often through features built into familiar business applications rather than standalone AI tools. As AI becomes more pervasive, organizations face a growing challenge in protecting sensitive information that may be exposed through both direct use and AI functionality operating behind the scenes.



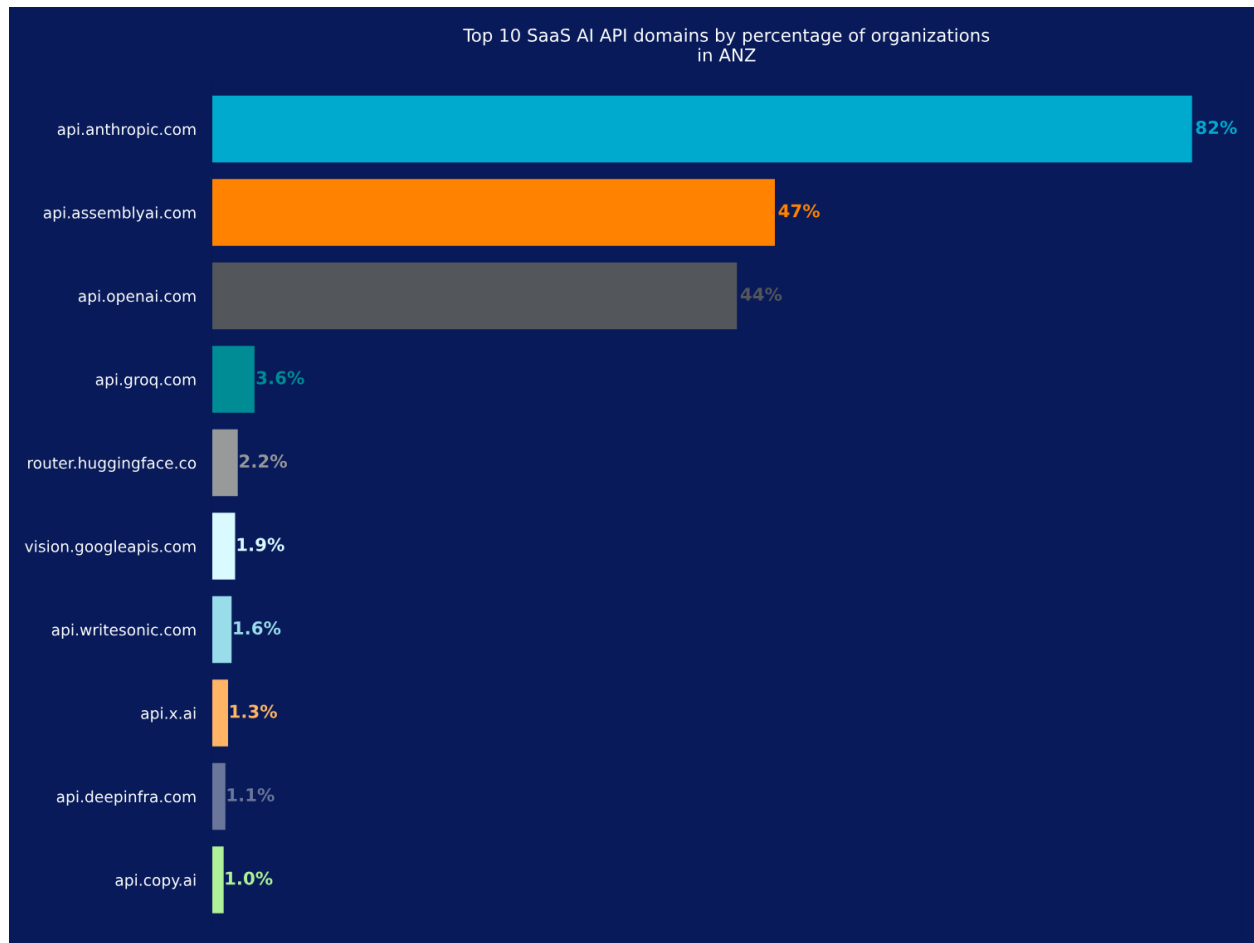
Agentic AI adoption

Rising use of AI APIs outside the browser

Even when AI agents and applications are deployed on-premises, they often rely on cloud-hosted models delivered through SaaS or enterprise AI platforms. Rather than interacting through a browser, these systems typically connect to AI models through dedicated APIs. For example, browser-based interactions with OpenAI generally occur through `chatgpt.com`, while internal applications and AI agents connect programmatically through `api.openai.com`.

Across ANZ, Anthropic's API (`api.anthropic.com`) is the most widely used AI API, with 82% of organizations connecting to it, well ahead of `api.assemblyai.com` at 47% and `api.openai.com` at 44%. The sizable lead mirrors the broader shift in AI adoption across the region, where Anthropic Claude has

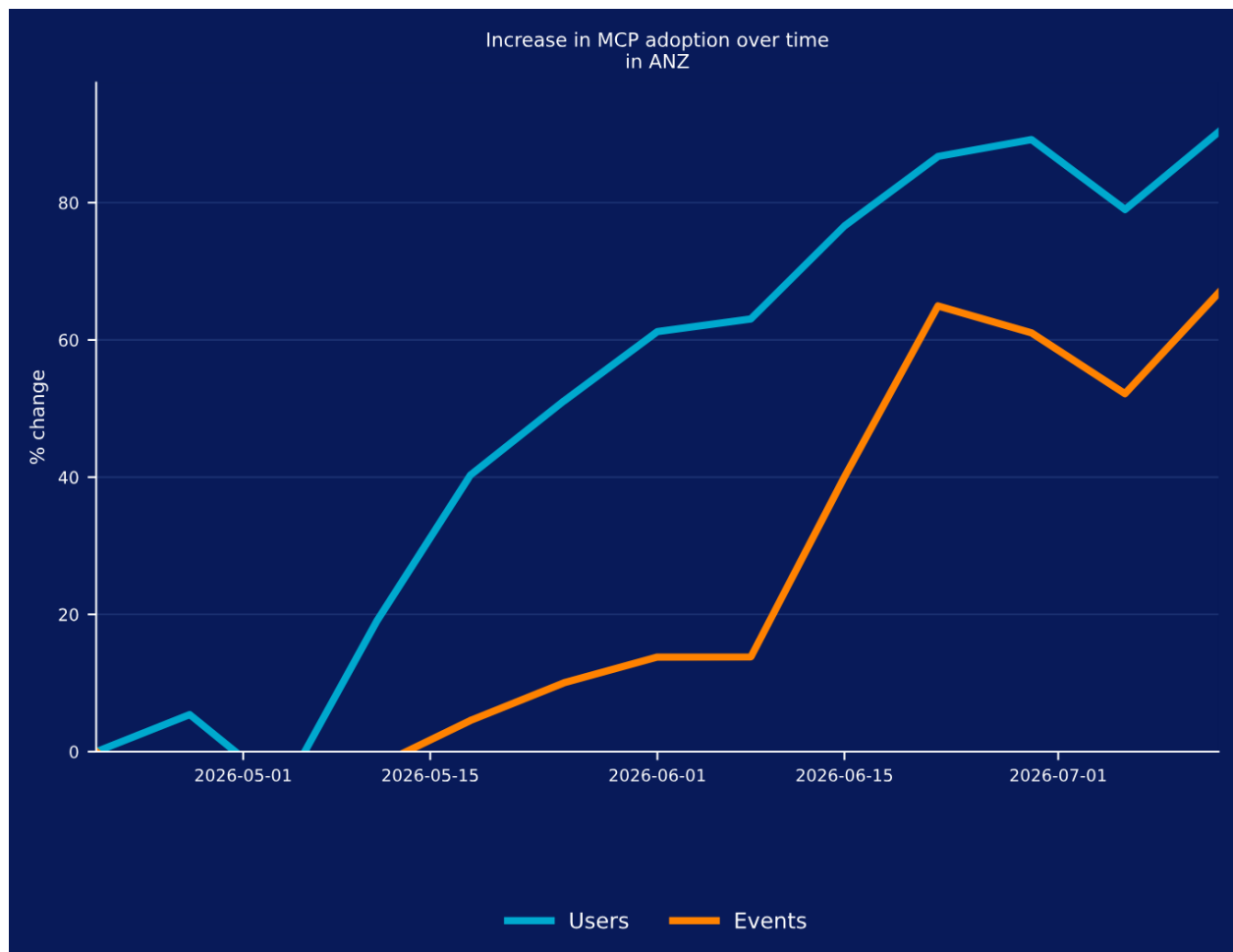
also overtaken ChatGPT as the most widely used AI application. This momentum is likely driven in part by the growing popularity of Claude Code and Anthropic models among developers, leading organizations to increasingly standardize on Anthropic for both end-user applications and API-driven AI integrations. As embedded AI services become more common across enterprise systems, securing and governing these API connections is becoming just as important as managing browser-based AI usage.



MCP: Rapidly increasing interconnectedness

MCP, the open-source standard for connecting AI models to data sources and tools beyond the model itself, is seeing rapid momentum across organizations in ANZ. Over the period, the number of agents interacting with remote MCP servers increased by 89%, while MCP-related events grew by 69%. This growth is significant because remote MCP connections create new pathways for enterprise data to flow between AI applications and external systems, increasing the importance of monitoring and governing

these interactions. These figures refer only to remote MCP usage, where AI agents connect to MCP servers hosted on the internet rather than on a local machine or within an organization's own network. Among the most common clients connecting to remote MCP servers are AI coding assistants such as Claude Code and Codex, highlighting the growing adoption of agentic AI in software development and other enterprise workflows.

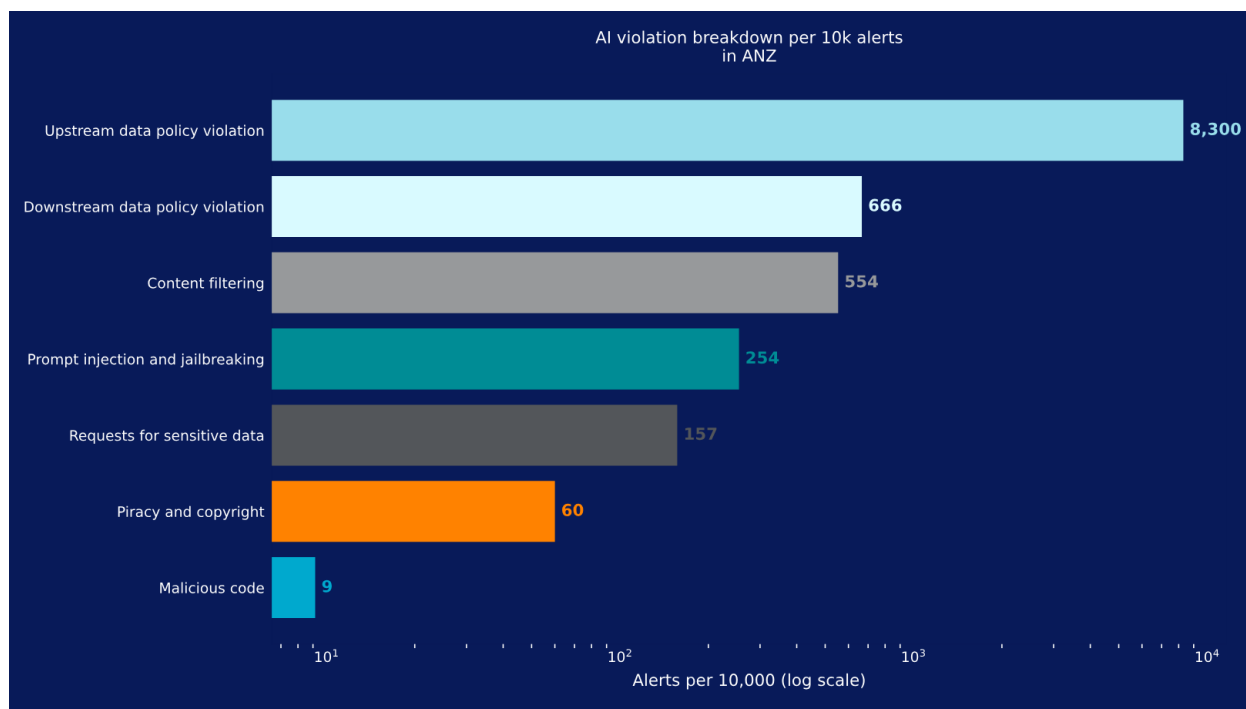


AI-adjacent threats

Categorizing AI Risks

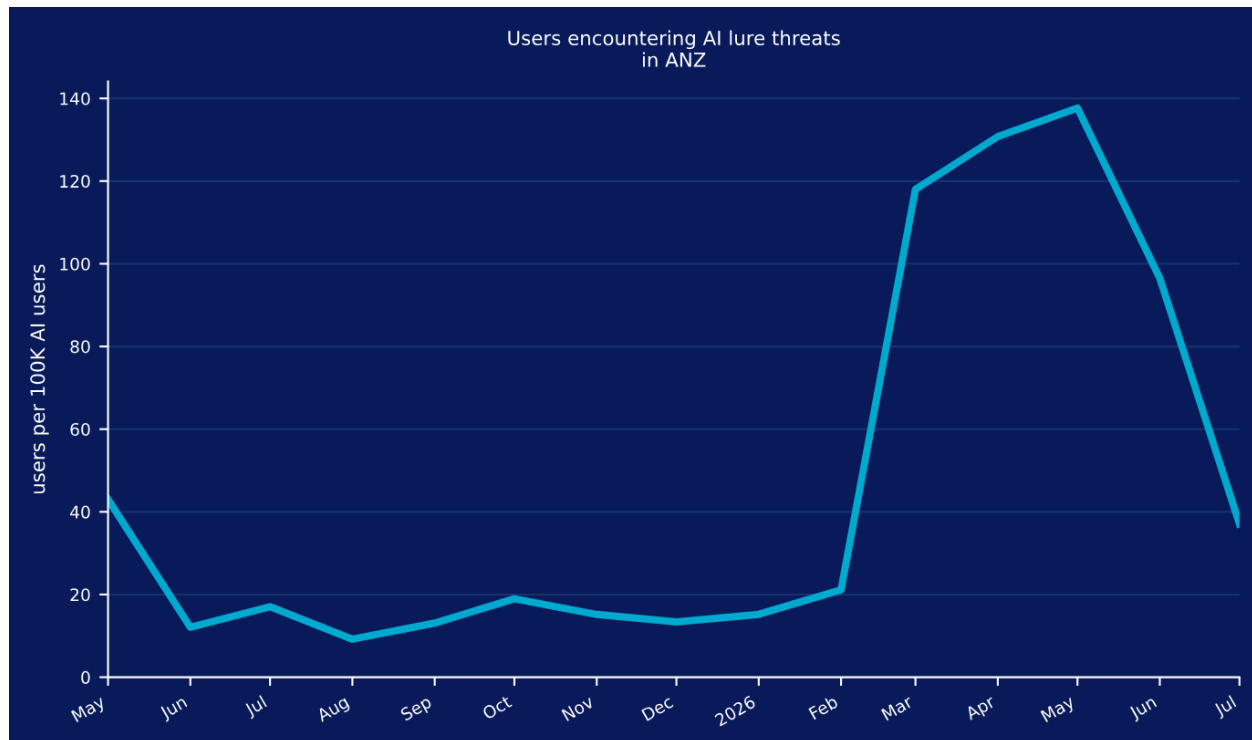
Effective AI visibility, governance, and protection begin with a clear understanding of the organization's risk landscape. The chart below ranks the most common AI-related risks based on how frequently they occur. Upstream data policy violations remain the most prevalent, affecting nearly all organizations, while

downstream data policy violations are also widespread, occurring in approximately 89% of organizations. Malware-related violations remain the least common, but they continue to represent some of the highest-impact risks due to their potential to compromise systems and data.



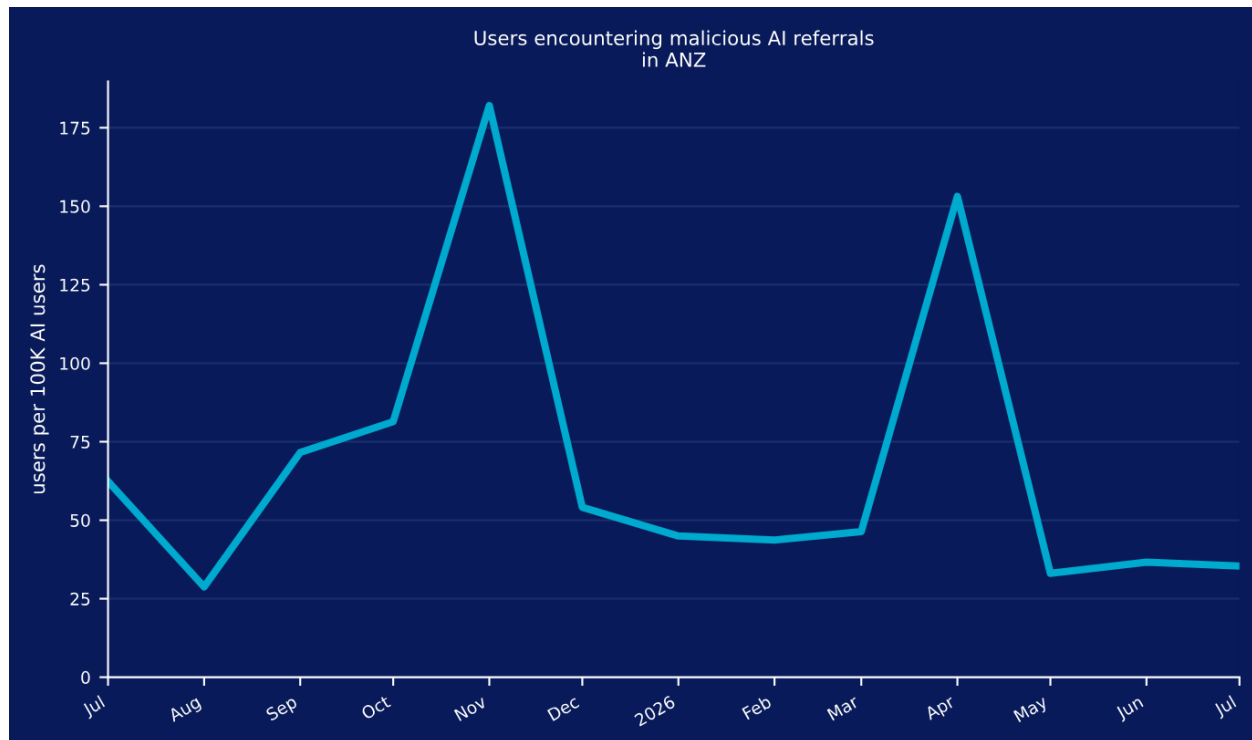
Malicious AI lures

A malicious AI lure attempts to trick victims into downloading malware or visiting a malicious website by impersonating trusted AI brands or tools. Across ANZ, AI lure activity remained relatively low and stable through most of the past year, before rising sharply in March 2026. The number of users encountering AI lures peaked in May at nearly 140 users per 100,000 before declining over the following two months, although activity remained well above earlier levels. Recent campaigns have used fake AI application [installers](#), [trojanized developer tools](#), and other AI-themed lures designed to exploit growing interest in AI. As organizations continue to adopt AI technologies and employees increasingly seek new AI tools, attackers are likely to keep refining these techniques, including targeting AI software supply chains through attacks similar to [Shai-Hulud](#) and publishing malicious packages that exploit AI-generated code recommendations.



Malicious AI links

A malicious AI link is a harmful link handed back by an AI app that a user clicks or an agent follows. They are the equivalent of malicious links attackers manage to reference on search engines such as Bing or Google, and that people click thinking they are legitimate domains. Those show up in results because the attacker used [SEO techniques to rank their malicious site highly](#), paid to slip a malicious ad into the results, or compromised otherwise benign infrastructure. In the AI world the technology has changed, but the playbook hasn't. Artificial intelligence engine optimization (AIEO) is emerging as the methods for getting content favored by AI models become better understood, and ads are starting to work their way into AI apps too. In ANZ, the rate at which users click such links has swung widely over the year, from around 26 to more than 175 per week per 100,000 users, and we expect the average to climb through the second half of 2026, despite the recent slowdown.

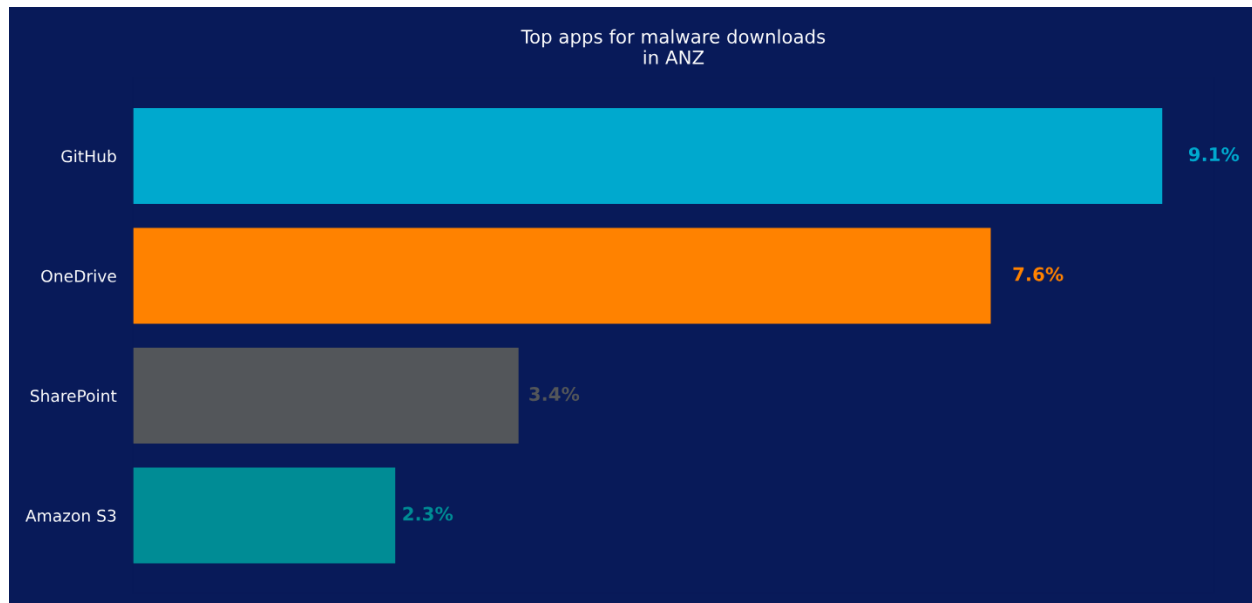


Malware downloads

Malware distribution via cloud apps

Attackers frequently exploit cloud platforms to distribute malware, leveraging user trust in well-known and legitimate cloud services. While providers actively remove malicious content, even short detection delays can allow successful infections and internal propagation.

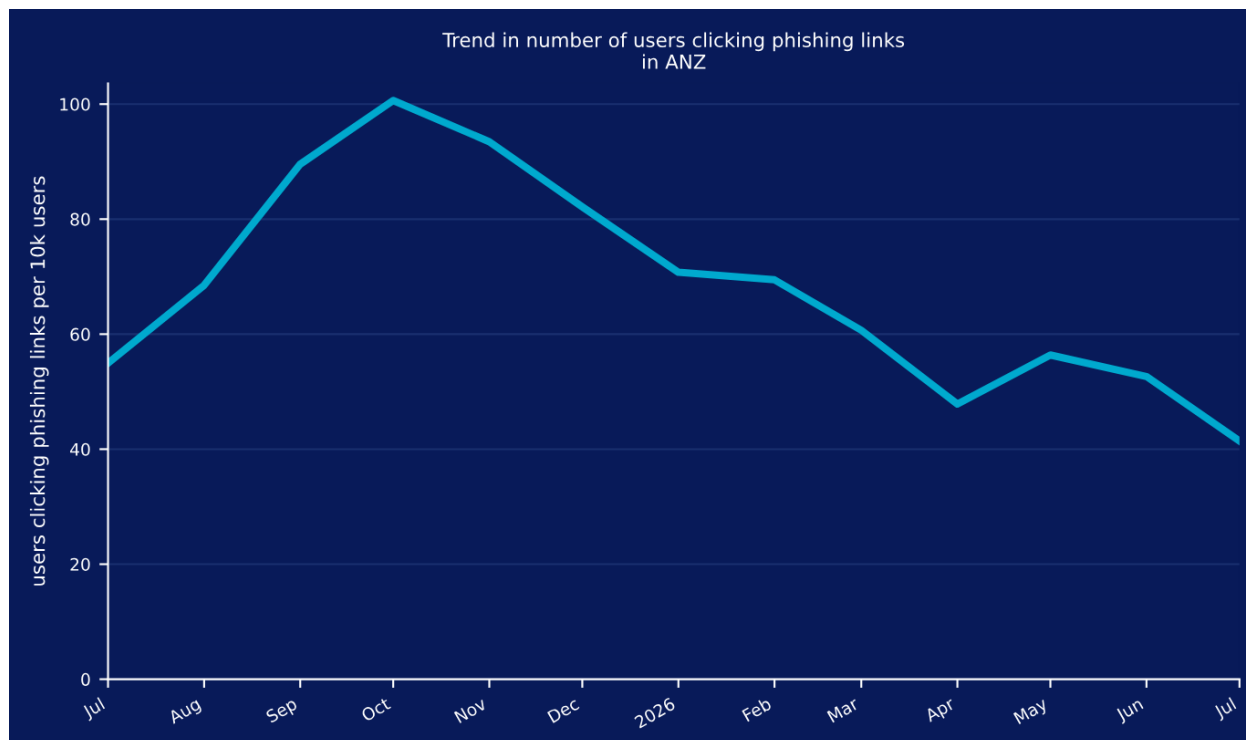
Across organizations in ANZ, GitHub and Microsoft OneDrive are among the most commonly abused platforms for malware distribution, affecting 9.1% and 7.6% of organizations respectively. This shows a broader shift in attacker behavior toward using legitimate cloud infrastructure rather than suspicious or unknown domains, and making malicious activity harder to detect within normal traffic patterns.



Phishing

Phishing remains one of the most successful techniques used by attackers, despite years of security awareness training and improvements in defensive controls. Attackers have steadily refined their methods, moving beyond simple credential-harvesting pages to techniques such as malicious OAuth applications, fake authentication portals, and reverse proxies that capture credentials and session tokens in real time.

Phishing susceptibility among users in ANZ has improved significantly over the past year. Since August 2025, the number of users clicking on phishing links has dropped by almost 60%, from 91 to 41 per 10,000 users. This trend suggests that organizations are becoming more effective at reducing phishing risk through a combination of stronger email security controls, increased user awareness, and broader adoption of phishing-resistant security measures.



Recommendations

With the growing use of AI tools, both managed and personal, and the misuse of personal cloud apps, it is essential to strengthen visibility, improve policies, and prioritize proactive defenses to protect your organization in this fast-changing threat landscape.

Based on the trends uncovered in this report, Netskope Threat Labs strongly encourages organizations across ANZ to take a fresh look at their overall security stance:

- Inspect all HTTP and HTTPS downloads, including all web and cloud traffic, to prevent malware from infiltrating your network. Netskope customers can configure their [Netskope One NG-SWG](#) with a threat protection policy that applies to downloads across all categories and all file types.
- Block access to apps that do not serve any legitimate business purpose or pose a disproportionate risk to the organization. A good starting point is a policy to allow reputable apps currently in use while blocking all others.
- Use [DLP](#) policies to detect potentially sensitive information, including source code, regulated data, passwords and keys, intellectual property, and encrypted data, being sent to personal app instances, AI apps, or other unauthorized locations.

- Use [Remote Browser Isolation \(RBI\)](#) technology to provide additional protection when visiting websites in categories that may pose a higher risk, such as newly observed or newly registered domains.
- Use [Netskope One AI Gateway](#) to gain visibility and control over AI applications and API interactions, helping secure data flows between users, applications, and LLMs.
- Deploy [Netskope One AI Guardrails](#) to enforce consistent protections against sensitive data exposure, unsafe prompts, and policy violations across managed and unmanaged AI environments.
- Use [Netskope One AI App Security](#) to discover sanctioned and unsanctioned AI applications, apply real-time controls, and enforce governance policies across personal and enterprise AI use.
- Leverage [Netskope One AI Analytics](#) to monitor AI adoption trends, user activities, and DLP incidents, facilitating organizations to understand better and reduce AI-related risk exposure.
- Consider [Netskope One AI Red Teaming](#) to actively identify vulnerabilities and misconfigurations in private AI deployments before they may be exploited in production environments.

Netskope Threat Labs

Staffed by the industry's foremost cloud threat and malware researchers, [Netskope Threat Labs](#) discovers, analyzes, and designs defenses against the latest cloud threats affecting enterprises. Our researchers are regular presenters and volunteers at top security conferences, including DEF CON, Black Hat, and RSA.

About this report

Netskope provides threat protection to millions of users globally. The information presented in this report is based on aggregate use data collected by the [Netskope One platform](#) for a subset of Netskope customers in ANZ.

The statistics in this report are based on the period from July 1, 2025, through July 15, 2026. Stats reflect attacker tactics, user behavior, and organization policy.

###