

AI as a Security Threat vs. Force Multiplier

We brought together a group of CISOs and senior security leaders to discuss:

- How AI is changing the threat landscape
- Where AI is strengthening security teams
- How teams are managing AI governance

STAYING AHEAD OF THE AI THREAT

Security leaders are balancing two competing questions:

- How can AI improve our defences?
- How do we control the risks it introduces?

Current tools remain forms of narrow intelligence rather than fully autonomous general intelligence.

The immediate challenge is not science fiction. It is the practical reality of sensitive data, uncontrolled usage, inconsistent governance and rapidly changing capabilities.

The biggest AI risk isn't the model. It's the lack of governance around how people are already using it.

Focus on clear ownership, practical guardrails and human accountability rather than trying to eliminate every risk.

IS ADOPTION OVERTAKING GOVERNANCE?

Many organisations are already dealing with shadow AI.

Teams have been using public AI tools for some time, often before formal policies, approved platforms or monitoring controls were introduced. This undoubtedly creates challenges around:

- Data sovereignty
- Sensitive information leaving the organisation
- Identity and access management
- Model and vendor assurance
- Ownership of AI-related decisions
- Visibility of unapproved tools and use cases

AI governance is often being handed to security teams by default. Debatably AI should be owned at an organisational or IT level, with security providing oversight, controls and challenge.

Create clear ownership before creating policy. Security should help define the guardrails, but should not become solely accountable for every AI decision made across the business.

TESTING IS TAKING A LAYERED APPROACH

Successful models combine continuous automated testing with targeted human-led investigation, treating AI as an additional layer of assurance, not a replacement for human judgment.

The biggest value is not simply identifying more issues. It is reducing noise, improving context and freeing experienced security professionals to focus on complex or high-impact work.

Human-accredited testing is still expected to remain important, particularly where regulators require recognised certification.

Teams are recommending:

- Annual accredited penetration testing for regulatory assurance
- Targeted testing throughout the year
- Agentic or automated testing within individual sprints and workflows

AI is already changing how security teams operate.

The organisations best placed to benefit will have clear ownership, practical guardrails and a realistic understanding of where human judgment still matters. Security leaders have an opportunity to move beyond being seen as blockers and instead help their organisations adopt AI in a way that is useful, proportionate and secure.

The human layer remains one of the biggest security risks. AI is now making shortcuts even easier, increasing the risk that users bypass controls in pursuit of speed or convenience.

Advice for security teams:

- Understand why users are taking shortcuts
- Provide approved tools that are easier to use
- Design safe processes around real working behaviour
- Create feedback loops between security and the wider business
- Make safe adoption more convenient than shadow usage

HUMAN ACCOUNTABILITY STILL MATTERS

AI is increasing automated decisions faster than many teams can review manually, leaving space for risks in:

- AI code that nobody fully understands
- Loss of business knowledge after reducing headcount
- Difficulty investigating incidents in AI systems
- Over-reliance on tools without understanding limitations

Automating repetitive decisions where the impact is reversible & keeping clear human accountability where decisions are consequential, complex or difficult to undo is key.