

Post event report

**AI SEC
SUMMIT**

AI Sec Summit

7th May 2026, London

Principal Sponsor



Strategic Sponsors



Education Seminar Sponsors



Branding Sponsors



Inside this report:

Sponsors
Key themes
Who attended?
Speakers
Agenda
Education Seminars



Speakers

Jonathan Armstrong, Partner, Punter Southall Law
Brett Ayres, CTO, Teneo
Max Berg, Principal Solutions Engineer, BeyondTrust
Bradley Boshier, Sales Engineer Manager, Varonis
Georges Bossert, Co-founder, Chief Technology and Product Officer, Sekoia.io
Donato Capitella, Principal Consultant, Reversesec
Robert Cooper, IT Security Engineering Lead, easyjet
James Derbyshire, VP, Strategic Partnerships, Harmonic
Adaora Ezennia, GRC Lead, THG PLC
Orlando Fernandez, Senior Technical Specialist at the Recovery, Resolution & Resilience team, Prudential Policy Directorate, Bank of England (BoE)
Bobby Filar, Head of AI, Sublime Security
Jon Harrison, Tech Lead, Local AI Division, Ministry for Housing, Communities & Local Government
Paul Jerram, Compliance and Responsible AI Officer, Keolis UK & Ireland
Yair Kler, Vice President, Security Architecture, DHL Group
Céleste Manenc, Senior Corporate Sales Engineer, CrowdStrike
Etay Maor, VP Threat Intelligence and Founding Member of Cato CTRL, Cato Networks
Tom McNamara, CEO, Atoro
Daniyal Naeem, Distinguished Engineer, Principal Security Authority – AI, BT
Ioan Nascu, GenAI Security Assurance specialist, Citi
Daniel Oxley, Senior Engineer, Doppel
Ahsan Qureshi, Managing Director – Cyber Risk Advisory and AI Security, Ankura
Ryan Rubin, Senior Managing Director – Cyber Security, Privacy and AI Security, Ankura
Manit Sahib, Ethical Hacker & Former Head of Penetration Testing & Red Teaming, Bank of England
Ali Shepherd, Director of Cyber & Operational Resilience (CISO), FCA
Natalia Shevchuk, SVP, AI Security Architect, Citi
James Sherlow, Global VP of Sales Engineering, Wallarm
Patrick Sullivan, VP Innovation and Strategy, A-LIGN
Sam Watts, Product Lead for AI Agent Security, Check Point Software Technologies
John Wood, Leader, Next-Gen Application Security, Contrast Security

Key themes

- Data protection, privacy & confidentiality leakage risks
- Secure AI model development & MLOps hardening
- AI-Augmented cyber-attacks
- Human-AI interaction & control boundaries
- Operational resilience & AI failure management
- Regulatory landscape, compliance & liability
- AI-driven identity security & insider threat detection
- AI-powered vulnerability discovery & code security
- AI anti-phishing & social engineering defences
- AI for supply-chain & dependency risk
- AI-enhanced SOC operations
- Intelligent threat detection & behavioural analytics

Who attended?



Agenda		
08:00	Breakfast networking and registration	
08:50	Chairman's welcome	
09:00	The challenge of securing AI on a global scale	
	Yair Kler, Vice President, Security Architecture, DHL Group - Learn how CISOs can enable responsible AI adoption by advancing innovation without resorting to a blanket 'no' while preserving strong security foundations - Understand how FOMO-driven AI adoption is pressuring enterprises into rapid, high-risk decisions while bypassing established best practices - Addressing the issue of emerging AI-specific threat classes such as prompt injection and why they create new risks that remain difficult to mitigate - Recognising how long-standing security challenges like secrets management and identity lifecycle governance are re-emerging with greater complexity in AI-driven environments	
09:20	Breaking the backbone: Measuring real LLM security failures in AI agents	
	Sam Watts, Product Lead for AI Agent Security, Check Point Software Technologies - LLMs undergo extensive safety and security testing before release. The problem is this testing doesn't help security and AI teams make practical decisions about which models to trust as the backbone of their AI applications and agents - Understand how your model choice changes your risk profile in measurable ways and the real things that matter in model supply chain risk (and the things that don't!) - Based off cutting-edge research by Lakera and the UK AI Security Institute	
09:40	To build taller walls, or stronger gates? Identity is a battlefield, but privileges are the gatehouse	
	Max Berg, Principal Solutions Engineer, BeyondTrust - How AI agents and automation are driving a surge in non-human identity creation and associated risks in these quietly gaining privileged access without human-level scrutiny - Why logging in is the new breaking in. How threat actors and agents take the shortest path to administrative control by exploiting identity relationships across directories, cloud platforms, and federated trust boundaries - Why privilege is the true point of convergence, and how just-in-time access, least privilege, and PAM help limit the blast radius - Visibility without privilege context is just noise. A privilege-centric lens turns identity alerts into actionable risk rather than just another dashboard	
10:00	Beyond human identity: Securing AI agents	
	Ioan Nascu, GenAI Security Assurance specialist, Citi - Agentic identity: Why AI agents require fundamentally different approaches to traditional Identity and Access Management (IAM) - Emerging threat landscape: The unique risks that arise when intelligent agents become threat actors or targets - Beyond human capacity: How oversight mechanisms work when AI agents outnumber, outpace, and outlast human administrators	
10:20	Education Seminars Session 1	
	Contrast Security We have built for a world that no longer exists John Wood, Leader, Next-Gen Application Security, Contrast Security	Varonis From prompts to permissions: The new data risk model for AI Bradley Boshier, Sales Engineer Manager, Varonis
		Wallarm Your AI deployments are outpacing your defences: How to protect AI systems in production James Sherlow, Global VP of Sales Engineering, Wallarm
11:00	Networking break	

Agenda

11:30	PANEL DISCUSSION Who owns AI risk? And how do we stay compliant?		
	Simon Brady, Event Moderator; Jonathan Armstrong, Partner, Punter Southall Law; Orlando Fernandez, Senior Technical Specialist at the Recovery, Resolution & Resilience team, Prudential Policy Directorate, Bank of England (BoE); Adaora Ezennia, GRC Lead, THG PLC; Paul Jerram, Compliance and Responsible AI Officer, Keolis UK & Ireland • In practice, who owns AI risk in your organisation – and is that ownership clearly defined at executive level? • How are you ensuring AI use across the business stays aligned with existing regulatory obligations? • What visibility do you have over third-party, embedded & shadow AI tools – and how does that impact your compliance posture? • If asked to evidence AI governance to the board or a regulator tomorrow, how confident would you be?		
11:50	AI vs AI: Navigating the new era of the cyber-battlefield		
	Céleste Manenc, Senior Corporate Sales Engineer, CrowdStrike • Artificial intelligence is changing the pace and scale of cyber-operations • Adversaries are using AI to accelerate reconnaissance, automate intrusion paths, and exploit weaknesses faster than traditional defences can respond • In this session, CrowdStrike shares frontline insight into how this shift is unfolding across the global threat landscape • We examine how threat actors are applying AI today and what effective, AI-native defence looks like in practice • The discussion focuses on practical decision-making, resilience, and how organisations can apply AI with discipline to stay ahead as adversaries continue to evolve		
12:10	Navigating AI risks: Practical risk management strategies for security and compliance teams		
	Patrick Sullivan, VP Innovation and Strategy, A-LIGN & Tom McNamara, CEO, Atoro • If your company has adopted AI tools and processes, you've also adopted the risks that come with them. Is your organisation prepared to manage these risks? • Join experts from leading auditor, A-LIGN, and security and compliance partner and customer, Atoro, to cut through the noise and learn practical, real-world strategies for managing your organisation's AI risk with ISO 42001 and ISO 27001 • Why acting now on AI governance is business-critical • The real-world case for ISO 42001 • Why combining certifications delivers maximum protection with greater audit efficiency		
12:30	Trust, then autonomy: A new framework for evaluating agentic AI in security		
	Bobby Filar, Head of AI, Sublime Security Over the past 18 months, the security industry has sprinted from 'AI-assisted' to 'autonomous'. At most large scale industry conferences you attend in 2026, the show floors are almost entirely dedicated to agentic AI, yet the evaluation rigor applied to those products has barely changed. This talk will cover: • How to identify the gap between the autonomy a vendor markets and the autonomy they actually deliver, and why that gap is more dangerous in security than anywhere else • A practical autonomy framework for security AI (the equivalent of SAE levels) and why advancing from Level 1 to Level 4 is a trust relationship between vendor and customer, built on transparency, explainability, and auditability, rather than a product roadmap • What rigorous AI evaluation actually looks like: benchmark design, test set construction, and why a vendor who can only show results (not methodology) is telling you something important • Six questions to ask any security vendor that claims to have autonomous AI, which cannot be answered with marketing speak		
12:50	Education Seminars Session 2		
	Reversec Every guardrail everywhere all at once Donato Capitella, Principal Consultant, Reversec	Sekoia.io Agentic SOC: Data rules Georges Bossert, Co-founder, Chief Technology and Product Officer, Sekoia.io	Teneo & Palo Alto Networks How to translate your AI policy into enforceable security controls Brett Ayres, CTO, Teneo
13:30	Lunch and networking		

Agenda

14:30	Designing trusted AI: Secure-by-default architectures and AI-enhanced SOC in practice		
	Daniyal Naeem , Distinguished Engineer, Principal Security Authority – AI, BT - How to design secure-by-default architectures for agentic AI systems, grounded in clear security policies and operational standards - A practical MCP secure reference architecture - Real-world use cases for AI-augmented SOC operations - Key risk considerations when operationalising AI in security environments		
14:50	Social engineering attack chain: A new standard for unified defence		
	Daniel Oxley , Senior Engineer, Doppel - Social engineering is no longer confined to isolated phishing attempts, but operates as a coordinated, AI-driven attack chain that spans email, SMS, voice, chat, and video, requiring a shift in how organisations understand and manage exposure - Gain insight into how attackers leverage AI to scale, personalise, and synchronise interactions across channels, creating campaigns that are more convincing and harder to detect using traditional controls - Explore how organisations can assess their own exposure to cross-channel, AI-enabled social engineering, and where gaps typically exist between perceived and actual risk - Develop an understanding of human risk management as a practical framework, moving beyond awareness programmes to measure and reduce risk at the human layer - Learn how to align security controls, operational processes, and business stakeholders to build a unified, intelligence-led defence model against human-targeted attacks		
15:10	Humans are the weakest link? Think again		
	Etay Maor , VP Threat Intelligence and Founding Member of Cato CTRL, Cato Networks - Challenge the current approach to AI agents security - Demonstrate risks and attacks against agents - Look into the dark web and criminal forums as to what threat actors are saying		
15:30	Education Seminars Session 3		
	Ankura Navigating AI risk: The security mythos, the ecosystem, and the anatomy of a breach Ryan Rubin , Senior Managing Director – Cyber Security, Privacy and AI Security, Ankura & Ahsan Qureshi , Managing Director – Cyber Risk Advisory and AI Security, Ankura	Check Point Software Technologies Securing the agentic age: A practical guide Sam Watts , Product Lead for AI Agent Security, Check Point Software Technologies	Harmonic The AI genie is out of the bottle James Derbyshire , VP, Strategic Partnerships, Harmonic
16:10	Networking break		
16:40	Invisible leaks: The hidden risks of chatting with AI		
	Manit Sahib , Ethical Hacker & Former Head of Penetration Testing & Red Teaming, Bank of England - AI privacy risks: How tools like ChatGPT, Claude, and Co-Pilot can end up knowing more about you than your best friend (and never forget a thing). The hidden dangers of casually sharing information with AI - When small details add up: Why a few 'harmless' details can combine to paint a full picture & how scattered information can reveal sensitive data without you realising - The myth of security: Why AI models aren't as secure as we might think & how attackers can trick them into spilling information - Simple, practical steps: For employees: how to keep personal and company data safe & for organisations: reducing AI-related risks before they grow		
17:00	PANEL DISCUSSION	Buying AI without buying risk	
	Simon Brady , Event Moderator; Robert Cooper , IT Security Engineering Lead, easyjet; Ali Shepherd , Director of Cyber & Operational Resilience (CISO), FCA; Natalia Shevchuk , SVP, AI Security Architect, Citi; Jon Harrison , Tech Lead, Local AI Division, Ministry for Housing, Communities & Local Government - When you're buying an AI product, what's the first security concern that comes to mind? - When an AI vendor says their product is 'secure', what do you actually want to hear from them? - What's the fastest red flag that makes you pause or stop an AI purchase? - What's one question every procurement team should ask before signing an AI contract? - If you could give one piece of advice to someone buying AI today, what would it be?		
17:30	Chairman's closing remarks	17:30	Drinks reception & networking
		18:30	End of conference

Education Seminars

Ankura

Navigating AI risk: The security mythos, the ecosystem, and the anatomy of a breach

Ryan Rubin, Senior Managing Director – Cyber Security, Privacy and AI Security, Ankura & **Ahsan Qureshi**, Managing Director – Cyber Risk Advisory and AI Security, Ankura

The rush to adopt AI has created a dangerous 'mythos' – the belief that applying standard frameworks and delegation of risk solely to the CISO are enough to keep organisation's safe. While there is no shortage of theoretical frameworks, organisations are struggling to secure the actual, complex AI ecosystem, at the rapid pace of adoption. This ecosystem spans infrastructure, applications, Model Context Protocols (MCP), and third-party packages.

Furthermore, Shadow AI evolves into autonomous agentic risks, the threat landscape is shifting under our feet. When this ecosystem fails fast, the resulting AI breach behaves, scales, and must be remediated completely differently than traditional cyber-incidents. Join us to explore practical, overarching risk strategies that move beyond theoretical checklists and prepare your entire business for the new reality of AI threats.

Attendees will learn:

- Shattering the AI mythos: Why standard frameworks and treating the CISO as the sole 'magic bullet' fall short.
- Securing AI: Identifying real-world vulnerabilities across the Ecosystem
- The evolution of shadow AI: Uncovering and managing hidden, autonomous agentic risks within your environment
- Why an AI breach is different: Understanding the unique anatomy, forensic challenges, and response strategies for AI-specific incidents
- Holistic risk management: Building a cross-functional defence that moves past 'AI governance' to practical resilience

Check Point Software Technologies

Securing the agentic age: A practical guide

Sam Watts, Product Lead for AI Agent Security, Check Point Software Technologies

We'll be sharing best practices and hard-won lessons from securing production AI use and deployments across fortune100 enterprises and globally important companies.

Attendees will learn:

- Get practical guidance in securing agent adoption across organisations at scale
- Learn what production AI security looks like in AI applications used by millions of global customers per day

Contrast Security

We have built for a world that no longer exists

John Wood, Leader, Next-Gen Application Security, Contrast Security

AI is accelerating both code creation and attack capability beyond the limits of traditional application security models.

Attendees will learn:

- AI is accelerating both code creation and attack capability beyond the limits of traditional application security models
- AI-driven development has collapsed deployment cycles while scanning programmes remain structurally slower than the code they are meant to assess
- AI-assisted attackers iterate at machine speed, rendering signature-based detection increasingly ineffective
- The result is a widening exposure gap: more unreviewed code in production, and more adaptive exploitation targeting it

Harmonic

The AI genie is out of the bottle

James Derbyshire, VP, Strategic Partnerships, Harmonic

The AI genie is out of the bottle and it's now acting autonomously. With 240 AI tools per company and agents embedded across every workflow, the workforce has outpaced IT governance entirely. James Derbyshire draws on data from 22 million enterprise AI prompts to explore how agentic AI is reshaping work, where the new security risks lie, and why the answer isn't to block the genie...it's to govern it.

Attendees will learn:

- Why agents change everything. They trigger, reason, act, and repeat without asking permission. That breaks every risk model built for the previous generation of AI tools.
- When agents go rogue. Deleted inboxes. A 13-hour AWS outage. A sandbox escape mining crypto. One in eight reported AI breaches now involves an autonomous agent.
- How to govern without blocking. Blanket bans failed. Learn what security-mature organisations are doing to let their workforce move fast without the exposure

Education Seminars	
Reversec Every guardrail everywhere all at once Donato Capitella, Principal Consultant, Reversec	This talk shares practical lessons learned from hands-on testing of real-world generative AI application use cases, focusing on how security failures emerge when LLMs are integrated into production systems. Attendees will learn: • The most common risks identified in LLM-enabled applications and the guardrails that are frequently missing • What 'good' looks like for LLM guardrails, and how those guardrails can be evaluated in practice • How to leverage guardrails in production as detection and response signals to defend against persistent attackers
Sekoia.io Agentic SOC: Data rules Georges Bossert, Co-founder, Chief Technology and Product Officer, Sekoia.io	Agentic SOC's are becoming inevitable, yet autonomy built on weak data leads to fragile decisions. In this session, Georges Bossert (CTPO) will demonstrate why 'data rules' and how to achieve reliable autonomy through a three-layer model: events, context, and cyber-threat intelligence (CTI). Attendees will learn: • How to apply a three-layer model (events, context, CTI) to build reliable autonomous SOC capabilities • How to design and implement TTP-guided runbooks for detection and response • The key guardrails required for safe autonomy, including traceability, confidence scoring, and stop conditions • How to deploy these capabilities while preserving data sovereignty, from cloud to on-prem and air-gapped environments without third-party exposure.
Teneo & Palo Alto Networks How to translate your AI policy into enforceable security controls Brett Ayres, CTO, Teneo	This session shows how to translate written AI usage guidelines into real, enforceable security outcomes using Palo Alto Networks' AI security capabilities. Attendees will learn: • Map AI policy intent directly to technical controls across users, apps, and data • Identify the right tools to detect shadow AI and policy violations • Enforce guardrails on GenAI access and Agents without killing productivity • Gain continuous visibility into AI-driven data exposure and misuse • Move from 'policy on paper' to measurable, auditable AI governance
Varonis From prompts to permissions: The new data risk model for AI Bradley Boshier, Sales Engineer Manager, Varonis	AI assistants and agents change how data is accessed, inferred, and exposed. New attack techniques exploit prompts, retrieved content, and permissions inside trusted systems, bypassing traditional controls. Attendees will learn: • Emerging AI-driven data risks • Why discovery alone isn't enough • How security teams can apply controls that scale with AI adoption • Why building a defensible AI security strategy is key for EU AI Act compliance
Wallarm Your AI deployments are outpacing your defences: How to protect AI systems in production James Sherlow, Global VP of Sales Engineering, Wallarm	When AI comes into your organisation, APIs multiply fast. Every agent, RAG pipeline, and third-party integration creates new endpoints that your security team often doesn't know exist. This session cuts through the theory to address the three operational challenges that matter most: finding the APIs you don't know about, protecting them without slowing down development, and reporting risk in terms leadership understands. Attendees will learn: • the most common areas of vulnerability, exploits and breaches • how to find the APIs that attackers can already see but you can't • ways to quantify the security debt your AI initiatives are accumulating