

BEYOND **DETECTION**: THE ROLE OF **AGENTIC AI** IN **CYBER** INVESTIGATION AND **RESOLUTION**

WHITEPAPER

LMNTRIX USA

19800 MacArthur Blvd,
Suite 850
Irvine, CA 92612
sales@lmntrix.com
888-388-1879

LMNTRIX AUSTRALIA

Level 25, 100 Mount street,
North Sydney 2060
sales@lmntrix.com
+61.288.805.198

LMNTRIX UK

Kemp House, 152 – 160
City Road, London, EC1V
2NX
sales@lmntrix.com
+44.808.164.9442

LMNTRIX SINGAPORE

60 Kaki Bukit Place, #05-19,
Eunos TechPark
sales@lmntrix.com
+65-3129-2639

LMNTRIX INDIA

VR Bengaluru, Level 5, ITPL Main
Rd, Devasandra Industrial Estate,
Bengaluru, Karnataka 560048, India
sales@lmntrix.com
+91-22-49712788

EXECUTIVE SUMMARY

Cybersecurity has long revolved around the problem of detection. Organizations deploy SIEM platforms, endpoint monitoring, anomaly detection engines, and behavioral analytics to identify suspicious activity. Detection has improved dramatically over the past decade, yet the flood of alerts that follow has created a new bottleneck. Security teams are inundated with signals, many of which are false positives or low-level anomalies. What is missing is not more detection, but the ability to answer the crucial question of what happens once a threat is detected.

Agentic artificial intelligence (AI) offers a path forward. Unlike traditional SOAR and rule-based systems, which rely on static playbooks and predefined triggers, agentic AI systems can investigate, reason, and even act in semi-autonomous or, in limited circumstances, act in an autonomous way. These systems can behave like Tier-1 analysts, taking an alert, gathering contextual information, developing hypotheses, testing them, and recommending or initiating resolution actions. The result is a closed loop that accelerates response, reduces workload, and improves consistency. This white paper explores how agentic AI can transform cyber investigation and resolution, drawing on recent academic research, industry case studies, and practitioner insights.



CONTENTS

EXECUTIVE SUMMARY	2
INTRODUCTION.....	4
KEY DEFINING AGENTIC AI IN THE CYBERSECURITY CONTEXT.....	5
HUMAN-AI TEAMING: REPLICATING THE TIER-1 ANALYST	6
CLOSING LOOPS IN PRACTICE.....	7
SOAR AND RULE-BASED SYSTEMS: STRENGTHS AND LIMITATIONS	8
RISK AND GOVERNANCE IN AGENTIC DEPLOYMENTS	10
IMPLEMENTATION PATHWAYS.....	11
A DEEP DIVE EXAMPLE	12
COMPARATIVE PERSPECTIVE: DETECTION, SOAR, AND AGENTIC AI	13
LMNTRIX'S AI ADOPTION AND USE AFTER DETECTION	14
SUMMARY OF HOW LMNTRIX ADDS VALUE DURING RESPONSE AND REMEDIATION WITH AI.....	16
ARTEMIS AND LISA IN MORE DETAIL.....	16
RECOMMENDATIONS AND ROADMAP	18
CONCLUSION.....	19



INTRODUCTION

Detection has become a commodity. Virtually every modern organization has multiple detection technologies in place, each producing alerts at high volume. The problem is that alerts alone do not solve the security challenge. SOC analysts often describe their day-to-day work as triage under fire: screens filled with new alerts, many of which demand attention but few of which contain enough context to know what to do next.

The result is alert fatigue, where critical signals risk being missed amidst the noise. Even when alerts are properly prioritized, investigating them is time-consuming. Analysts must gather logs, endpoint telemetry, user history, and external intelligence before they can assess whether an alert represents a real threat. This investigative delay increases attacker dwell time, giving adversaries opportunity to move laterally, escalate privileges, and exfiltrate data.

The cost is not just in risk but in resources. Security teams spend disproportionate amounts of time on low-value triage work, rather than strategic activities like proactive threat hunting or red-team simulation. Organizations therefore find themselves paying heavily for detection systems without seeing proportional improvement in outcomes.

The urgency is clear. Without a mechanism to rapidly transform alerts into action, detection loses much of its value. What security leaders need is not more alerts, but systems that answer definitively what to do next when alerts appear.



DEFINING AGENTIC AI IN THE CYBERSECURITY

CONTEXT

Agentic AI is a subset of artificial intelligence characterized by autonomy, adaptability, and the ability to execute multi-step reasoning. In cybersecurity, it refers to AI-driven systems that can take in an alert, orchestrate investigative steps across multiple data sources, form hypotheses about what is happening, and initiate or recommend a response.

Unlike simple automation, agentic AI is not bound by rigid if-then rules¹. It can dynamically decide which tools to call, how to interpret ambiguous signals, and when to escalate to human oversight. This adaptability makes it far more resilient in the face of novel or sophisticated attacks, where predefined playbooks may fail.

Agentic AI also differs from conventional AI classification models. While detection models might classify an event as benign or malicious, agentic AI is tasked with acting on that classification. It coordinates across systems, applies reasoning, and follows through until the incident is resolved or escalated. It is this focus on decision-making and action that makes the technology transformative.

¹ Nir Khestri, "Transforming cybersecurity with agentic AI to combat emerging cyber threats", Telecommunications Policy Volume 49, Issue 6, July 2025, <https://www.sciencedirect.com/science/article/pii/S0308596125000734>

HUMAN-AI TEAMING: REPLICATING THE TIER-1 ANALYST

One of the most compelling aspects of agentic AI is its potential to operate like a Tier-1 analyst. Recent research describes models in which agents assume investigative roles typically handled by junior analysts². These roles include triaging alerts, enriching them with context, generating hypotheses, and recommending responses. Human analysts retain oversight, stepping in at decision points that carry greater risk or ambiguity.

This co-teaming model demonstrates several benefits. First, it increases throughput, allowing more alerts to be processed without overwhelming human staff. Second, it introduces consistency, reducing the variability in how different analysts handle similar cases. Third, it frees senior analysts to focus on complex threats rather than repetitive triage.

The governance dimension is equally important. A 2025 working paper introduced the Model-Control-Policy (MCP) framework, which sets boundaries for agentic behavior³. The *Model* represents the underlying AI logic, the *Control* provides operational guardrails to prevent unintended action, and the *Policy* defines organizational rules that determine when agents act autonomously and when they defer to humans. This layered governance allows organizations to experiment with agentic autonomy while maintaining safety and accountability.

A further study on the evaluation of autonomous cyber defence agents showed the real-world impact of such systems. In a simulated enterprise environment, an agent enriched alerts, developed hypotheses about lateral movement, and executed containment actions under human validation⁴. The result was a reduction in attacker dwell time by nearly half, alongside a significant drop in analyst workload. These findings suggest that agents can indeed function as effective Tier-1 equivalents, particularly when paired with thoughtful governance.

² Massimiliano Albanese, Daniel Lende, Kevin Lybarger, Xinming Ou, "Towards AI-Driven Human-Machine Co-Teaming for Adaptive and Agile Cyber Security Operation Centers", arxiv, 09 May 2025, <https://arxiv.org/html/2505.06394v1>

³ August Moore, Ant Burke, Myles Foley, Anna Knack, Chris Hicks, Vasilios Mavroudis, "A Fundamental Research Plan for Autonomous Cyber Defence", CETAS, 13 May 2025, <https://cetas.turing.ac.uk/publications/fundamental-research-plan-autonomous-cyber-defence>

⁴ Johannes Loevenich, Erik Adler, Tobias Hürten, Roberto Rigolin F. Lopes, "Design and evaluation of an Autonomous Cyber Defence agent using DRL and an augmented LLM", Computer Networks, Volume 262, May 2025, <https://www.sciencedirect.com/science/article/abs/pii/S1389128625001306>

CLOSING LOOPS IN PRACTICE

Several recent deployments and case studies illustrate how automation is beginning to close investigative loops in real environments.

In one testbed, researchers evaluated an Autonomous Cyber Defence (ACD) agent capable of responding to alerts with forensic enrichment, hypothesis generation, and containment. Human approval was required for high-impact actions, but even with these safeguards, the agent reduced response time dramatically and improved root cause visibility⁵.

Vendor case studies offer further insights. In a vendor's analysis of security-focused agents highlights use cases such as automated triage, enrichment of alerts with contextual data, and initiation of playbooks⁶. Another vendor similarly catalogues seven use cases, including autonomous phishing investigation and automated execution of remediation tasks⁷. These vendor perspectives, though promotional, show how agentic workflows are being incorporated into SOCs incrementally, often beginning with semi-autonomous functions and gradually extending autonomy as confidence grows.

In a press release titled "Internet of Agents" a report provides a much broader view, demonstrating how agentic capabilities are not only being deployed defensively but are also appearing in attacker toolkits⁸. This dual-use reality underscores the need for defenders to embrace agentic systems if they are to keep pace with adversaries. Failing to adopt such tools risks leaving defenders outmatched in speed and adaptability.

⁵ Johannes Loevenich, Erik Adler, Tobias Hürten, Roberto Rigolin F. Lopes, "Design and evaluation of an Autonomous Cyber Defence agent using DRL and an augmented LLM", Computer Networks, Volume 262, May 2025, <https://www.sciencedirect.com/science/article/abs/pii/S1389128625001306>

⁶ Exabeam, "Security-Focused AI Agents: Benefits, Capabilities and Use Cases", Exabeam, 24 September 2025, <https://www.exabeam.com/explainers/agentic-ai/security-focused-ai-agents-benefits-capabilities-use-cases/>

⁷ Charlie Klein, "7 Use Cases for AI Agents in Cybersecurity", Jit.io, 2 June 2025, <https://www.jit.io/resources/devsecops/7-use-cases-for-ai-agents-in-cybersecurity>

⁸ Radware, "Autonomous AI Agents Expand Attack Surface: Key Insights from Radware's "Internet of Agents" Report", Radware, 17 September 2024, <https://kbi.media/press-release/autonomous-ai-agents-expand-attack-surface-key-insights-from-radwares-internet-of-agents-report/>

SOAR AND RULE-BASED SYSTEMS: STRENGTHS AND LIMITATIONS

To understand the promise of agentic AI, it is important to contrast it with existing automation systems such as those seen as Security Orchestration, Automation, and Response (SOAR).

SOAR platforms emerged as a way to automate responses to alerts by chaining predefined actions together. If an alert met certain criteria, the SOAR system would execute a playbook such as an isolating an endpoint, blocking an IP, or opening a ticket. While useful, this model has several limitations. Playbooks are brittle, requiring constant maintenance as environments evolve. They lack adaptability when faced with novel or ambiguous threats. And because their logic is static, they cannot generate or test hypotheses in the way human analysts do.

Agentic AI, by contrast, offers reasoning and adaptivity. It can re-evaluate midstream, shifting investigative direction as new information emerges. It can coordinate tools dynamically, pulling in logs, endpoint data, or external threat intelligence as needed. Furthermore, the AI can also learn from feedback, refining its actions over time rather than relying solely on manually updated playbooks.

The table below summarizes these contrasts:

Feature	SOAR / Rule-Based	Agentic AI
Adaptivity / recursion	Fixed playbooks; limited flexibility	Re-evaluates midstream, adjusts based on discoveries
Reasoning / hypothesis generation	Minimal	Generates and tests hypotheses like a human analyst
Tool coordination	Predefined integrations	Dynamic orchestration across diverse systems
Human oversight	Manual gates; high false positives	Policy-driven oversight; semi-autonomous or autonomous
Learning and feedback	Manual updates required	Continuous learning and adaptation

Practitioner commentary reinforces this distinction. Vendors have argued that many organizations under-utilize SOAR because of its maintenance burden and lack of adaptability. Further, vendors like promoting Agentic AI adoption emphasize that agentic systems surpass SOAR by enabling recursive decision-making and dynamic tool-chaining.



RISK AND GOVERNANCE IN AGENTIC DEPLOYMENTS

The promise of agentic AI must be balanced against its risks. Allowing software agents to act autonomously in sensitive environments introduces the possibility of unintended consequences.

One risk is erroneous action. An agent that isolates the wrong endpoint or resets the wrong credentials could disrupt business operations. Overreach is another concern: without carefully defined policies, an agent might act in areas beyond its intended scope. Equally, pressing is the issue of explainability. If an agent takes an action without producing a clear rationale, post-incident analysis and accountability become difficult.

Adversarial misuse is an emerging threat as well. Attackers might attempt to manipulate or impersonate defensive agents, causing confusion or inducing harmful actions. Finally, regulatory and legal considerations cannot be ignored. In many industries, actions such as data deletion or account modification are tightly governed. Agents acting without human oversight may breach compliance obligations.

Governance frameworks help mitigate these risks. The MCP model provides a structured way to define agent behavior through layered controls. Human-in-the-loop models allow organizations to calibrate autonomy gradually, beginning with semi-autonomous functions and increasing autonomy only after performance is validated. Continuous monitoring, red-team testing, and drift detection ensure that models remain reliable over time.

Transparency and auditability are essential. Agents should produce structured justifications for their actions, enabling analysts and auditors to trace decisions. Post-mortems should assess not only the incident but the agent's performance. In this way, organizations can build confidence while maintaining accountability⁹.

⁹ Phaedra Boinodiris, Jon Parker, "The evolving ethics and governance landscape of agentic AI", IBM, 21 March 2025, <https://www.ibm.com/think/insights/ethics-governance-agentic-ai>

IMPLEMENTATION PATHWAYS

Adopting agentic AI is not an all-or-nothing proposition. Organizations can follow a staged approach that balances ambition with caution.

The first step is identifying suitable use cases. High-volume, repetitive alert categories such as phishing or endpoint malware are strong candidates. These domains are well understood, carry moderate risk, and benefit from rapid triage.

Once use cases are defined, organizations should begin with hybrid modes. In these setups, agents gather data, enrich alerts, and suggest actions, but human analysts retain final approval. This model provides immediate value while limiting risk. Over time, as confidence in agent performance grows, autonomy can be increased.

Integration is a critical success factor. Agentic AI relies on access to diverse data sources and tools, from SIEM logs to endpoint telemetry and threat intelligence feeds. Organizations must ensure these integrations are robust.

Metrics are essential for evaluating success. Key measures include mean time to investigation, false positive rates, dwell time reduction, and analyst workload savings. Regular reporting builds confidence among stakeholders and informs iterative improvement.

Risk mitigation should be planned from the outset. Organizations should implement rollback mechanisms, manual overrides, and sandbox testing. They should also ensure comprehensive documentation and audit trails to meet compliance and governance requirements.



A DEEP DIVE EXAMPLE

Consider a mid-sized financial institution facing overwhelming alert volumes, particularly around phishing and anomalous logins. Traditionally, analysts spent hours enriching these alerts with contextual information before escalating or resolving them.

By deploying an agentic AI system, the institution restructured its workflow. The agent automatically enriched phishing alerts with sender history, user behavior patterns, and endpoint telemetry. It generated hypotheses about whether an account was compromised, tested these hypotheses against log data, and presented recommended actions such as quarantining emails or prompting a password reset.

Initially, all recommendations required human approval. Over time, as accuracy was demonstrated, the organization permitted the agent to autonomously quarantine low-risk phishing emails while continuing to escalate high-impact actions to humans.

The results were measurable. Time to resolution dropped by half, analyst workload decreased significantly, and the quality of investigations improved. Analysts reported greater confidence in root cause identification, and senior staff were freed to focus on advanced threat hunting.

This case illustrates the incremental, governed adoption of agentic AI, showing how organizations can realize benefits without taking unacceptable risks.



COMPARATIVE PERSPECTIVE: DETECTION, SOAR, AND AGENTIC AI

A side-by-side comparison of detection, SOAR, and agentic AI illustrates the evolution of cybersecurity operations.

Dimension	Pure Detection	SOAR (Rule-based Automation)	Agentic AI
Speed of response	Low	Moderate	High
Contextual understanding	Minimal	Limited	Strong
Adaptability	Weak	Weak	Strong
Human oversight	Heavy	Moderate	Adjustable
Risk of unintended action	Low	Moderate	Higher, but manageable
Maintenance burden	Moderate	High	Moderate to high, but improving

This comparison highlights that agentic AI is not without challenges, but its adaptability and capacity for reasoning offer advantages that neither detection alone nor SOAR can deliver.



LMNTRIX'S AI ADOPTION AND USE AFTER DETECTION

In the response and remediation phase of threat management, LMNTRIX's MXDR and XDR platforms leverage automation and artificial intelligence to accelerate outcomes. At the core is Artemis, the embedded AI analyst that can best be described as an agentic AI that hunts, investigates, and responds. Once a threat is detected, Artemis can act autonomously or semi-autonomously across endpoints, cloud, identity, mobile, and operational technology. These actions are reinforced by automation playbooks that execute containment, isolation, rollback, and forensic steps without requiring full reimaging, allowing the system to surgically remove malicious artifacts or reverse harmful changes.

Automation also drives containment and remediation. The platform can isolate compromised endpoints, block malicious processes, quarantine traffic, roll back system changes, and execute remediation scripts across multiple devices. All actions are coordinated from a unified console and can be triggered automatically, reducing reliance on manual workflows and ensuring faster, consistent enforcement across environments.

A common challenge after detection is separating real threats from noise. LMNTRIX developed its system and technology stack to reduce false positives by about 95 percent through machine learning and filtering logic. Alerts are enriched with contextual data, scored by severity, and prioritized automatically, enabling analysts to focus on genuine risks. As telemetry grows, detection and triage improve over time, mitigating alert fatigue and streamlining operations.

Following containment, root cause analysis and forensics clarify how incidents unfold. LMNTRIX employs packet capture, session reconstruction, and retrospective analysis to track attacker behavior across time. Automated root cause analysis consolidates findings into a clear sequence of events, correlating activity across endpoints, networks, and cloud environments to create a complete breach timeline.

To strengthen defenses, LMNTRIX uses Automated Attack Validation, which simulates real-world tactics such as lateral movement or data exfiltration. These exercises test whether detection and response mechanisms work as intended, exposing gaps in automation or playbooks. Results feed back into the system, continually refining defenses.

The platform also incorporates deception and disruption. By deploying decoys, honeypots, and breadcrumbs, it misleads attackers, gathers intelligence, and redirects malicious activity into safe zones. This provides defenders more time to remediate the real environment while denying adversaries straightforward access.

Human expertise remains central. LMNTRIX positions its AI as a force multiplier, handling repetitive triage and orchestration while analysts focus on policy, complex cases, and playbook refinement. Dashboards, enriched context, and attack chain visualizations support faster, more informed decisions.

LMNTRIX underlines these capabilities with performance metrics, reporting a mean time to detect under one minute and mean time to remediate under 30 minutes for most incidents. By combining agentic AI, automation, deception, and human collaboration, the platform accelerates recovery and ensures resilience against evolving threats



SUMMARY OF HOW LMNTRIX ADDS **VALUE DURING**

RESPONSE AND REMEDIATION WITH AI

Putting it all together, here's how LMNTRIX's post-detection phase is enhanced by AI / automation:

Capability	Role / Benefit
Automated playbooks & orchestration	AI triggers containment, rollback, isolation, and remediation actions automatically or semi-automatically
Alert triage & prioritization	Reduces noise, highlights real threats, accelerates decisioning
Root cause & forensic correlation	AI helps reconstruct attack chains and attribute root causes across domains
Adversary emulation / validation	Exercises the remediation logic to uncover gaps in defenses
Deception / traps	Diverts attacker progression, gathers intelligence, buys time to respond
Human + AI partnership	AI handles the heavy lifting; analysts intervene on tricky or strategic decisions
Faster time metrics	Enables sub-minute detection and 30-minute or lower remediation for many incidents

ARTEMIS AND LISA **IN MORE DETAIL**

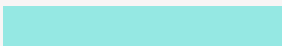
Artemis and LISA are two complementary components within the LMNTRIX XDR platform that work together to deliver faster and more effective detection, remediation, and response. Artemis is the AI-driven detection and response engine, constantly analyzing telemetry from endpoints, networks, identities, and cloud services to identify anomalies and threats in real time. It excels at speed, rapidly triaging, prioritizing, and even automating responses such as isolating infected systems or disabling compromised accounts.

LISA, on the other hand, enriches those detections with intelligence and context. It aggregates threat data from hundreds of sources, validates alerts, and provides deeper insight into incidents through forensics, identity monitoring, and threat correlation. This reduces false positives, supports investigators with detailed context, and ensures response actions are both accurate and proportionate. While Artemis drives automation, LISA ensures that decisions are guided by intelligence and that analysts have the necessary visibility into what's happening.

Together, they create a balance of speed and precision: Artemis accelerates detection and containment, while LISA ensures depth, context, and long-term remediation. This synergy enables security teams to cut through noise, respond quickly to genuine threats, and learn from incidents to strengthen defenses over time. With human analysts still in the loop for oversight, the Artemis-LISA combination provides both automated efficiency and the contextual intelligence needed for resilient cybersecurity operations.



RECOMMENDATIONS AND ROADMAP



Organizations considering agentic AI should begin by assessing readiness. Do they have the telemetry and data integration required for agents to reason effectively? If not, improving visibility is the first priority.

Once readiness is established, organizations should define narrow use cases with clear boundaries. Semi-autonomous modes should be the default in early pilots, ensuring that human oversight remains strong. Metrics must be carefully tracked to evaluate performance.

Over time, autonomy can be expanded. Organizations should align governance with regulatory requirements and risk appetite, ensuring transparency and auditability throughout. By adopting an iterative approach, organizations can capture the benefits of agentic AI while avoiding the pitfalls of overreach.



CONCLUSION

Detection is essential, but on its own it is insufficient. The modern threat landscape demands speed, adaptability, and closed investigative loops. Agentic AI provides these capabilities, offering systems that can think and act like analysts, enrich alerts, develop hypotheses, and execute responses under governance.

The comparison with SOAR systems makes clear that static, rule-based automation cannot keep pace with evolving threats. The evidence from research studies and vendor deployments demonstrates that agentic AI is already delivering real benefits, from reduced dwell time to improved analyst productivity.

Risks remain, but with frameworks like the MCP model, human-AI teaming, and robust governance, these risks can be mitigated. For organizations willing to adopt agentic AI carefully and incrementally, the reward is a SOC that can move beyond detection to investigation and resolution.

In a security landscape defined by speed and complexity, answering the question *"what now?"* is no longer optional. Agentic AI offers an answer to what can be done after detection.

