

WHITE PAPER

Reengineering the SOC: A Roadmap to AI-Enhanced Cyber Defense

Written by **Christopher Crowley**
July 2026



Executive Summary

Artificial intelligence (AI) and machine learning (ML) have compelled companies to take another look at how they're working. This is the reality of business operations today. The fundamental goals of business haven't changed much since before GenAI captured the public imagination, but GenAI has exponentially increased the internal pressure to perform and optimize. This is true for every component of the organization, especially cybersecurity operations. Why? The perception that disruption is more readily available to competitors.

The push to move to the next generation typically involves industry best practices review, organizational introspection, and envisioning the bold future. Implementing those future demands requires change. The risk is massive churn with little true progress accomplished.

This paper covers substantial ground. The core message is worth stating upfront, before the supporting analysis unfolds:

- Using AI doesn't require understanding how it works, but positioning your organization for future success in a tumultuous AI marketplace does.
- Organizations depend on managed security service providers (MSSP). Selecting the right one requires self-awareness, savvy shopping and contracting, and effective integration of the partner into your analytical cybersecurity workflow.
- AI helps cybersecurity operations with visibility, observability, and authorization. However, it doesn't do so consistently in all circumstances nor in the same way for each organization.
- AI tools enhance communication, automation, analysis, and detection engineering.
- AI maturity will be delivered incrementally, and will require organizational effort to integrate this new and rapidly transforming technology.
- Organization economic constraints and inflexibility will diminish benefits of AI.
- Internal staff competent in cyber-first principles, AI for cyber, engineering, and MSSP/outsourced engagement will be invaluable. Without them, an organization will be inept in using its IT resources; and at heightened risk of cyber compromise.

The white paper provides current state, near-term, and long-term perspectives on the use of AI in cybersecurity and how AI will transform cybersecurity.

Introduction

AI/ML has forced every part of the organization to take a hard look at how it operates, especially cybersecurity operations. The same tools that promise to enhance detection, analysis, and response arrive faster than any security team can properly vet them. They come from vendors whose own security practices are visibly imperfect into environments already struggling to keep pace. The opportunity is real and so is the risk of getting it wrong.

This paper moves from foundational knowledge to current to near-term and then long-term guidance. It is a roadmap for getting it right, from understanding what AI/ML actually does, through the near-term operational adjustments that position teams for effective integration to the longer-horizon model that serious security operations should be building toward. That model is not simply a better-equipped SOC, but rather a different operational paradigm where cybersecurity converges back into IT operations as a mature discipline, with AI as the force multiplier that makes that integration possible at scale (see Figure 1).

Keep in mind that AI is basically math, specifically predictions based on patterns in training data. A model trained on the wrong data or applied to problems requiring business context it cannot access will produce confident-sounding predictions that are wrong in ways that are hard to detect and potentially catastrophic when acted upon.

We are at an inflection point. The decisions security leaders make now about tools, staff, and the relationship between AI capability and human judgment will define their organizations' security posture for the next decade. This paper is intended to help inform those decisions.

CURRENT STATE

Reactive, Tool-Driven

- Shadow IT AI adoption
- Vendor-pushed integrations
- Detection gaps
- Siloed MSSP relationships

NEAR TERM

Engineered, Integrated

- Hunting-to-detection pipeline
- MCP and RAG implementations
- Structured MSSP collaboration
- Analyst-developer-engineer roles

LONG TERM

Converged, Resilient

- Security embedded in IT ops
- AI at operational scale
- Distributed security expertise
- Knowledgeable staff as force multipliers

Figure 1. The SOC Evolution Roadmap

Foundational Knowledge

Strategic deployment of AI/ML tools requires more than confidence in their outputs. It requires understanding the mechanism well enough to anticipate where predictions break down and why. What follows is the working knowledge security practitioners need to use these tools without being misled by them.

How AI/ML Works

Deep learning is a technique for prediction. To understand why AI tools have become de facto overnight, we don't need to understand how they work. We can just see their value. To differentiate between what we can see now and what the strategic implementation of these tools should be for the future, we need our own predictive capability. That means understanding the mechanism well enough to know its shortcomings.

Consider asking Google Gemini, "What do you think ChatGPT would say about the development of AI tools over the next five years, and would your analysis and predictions be different?" Gemini would summarize, "If I had to summarize the difference: ChatGPT would likely predict a future defined by how deeply AI can think (reasoning and agents), whereas I predict a future defined by how broadly AI can perceive (multimodality and physical-world integration)."

How does it arrive at that? Without going into lots of underlying math, it has meaningfully grouped the words provided and developed a series of words that describe the concept being asked about. It turns words into numerical relationships, uses those relationships to arrive at a set of other relationships that address the inquiry, then translates those numbers back into words. It has no idea what any of the words mean. It has a massive trove of relationships expressed as numbers, and it echoes those numbers as language. That is a deliberate oversimplification, but it is accurate enough to make the critical point.

Trusting AI/ML output is not the same as understanding it, and only the latter allows a practitioner to anticipate where the tool will fail.

The Failed Login

Large language models (LLMs) are extraordinarily complex. To explain the predictive mechanism more clearly, consider a simpler scenario of a series of failed login attempts followed by a successful login on a single account. This could be a user who mistyped a password or an eventually successful password-guessing attack by an adversary.

A security analyst investigating this, without asking the user, would consider a range of factors including historical data about IP addresses commonly associated with the account, the timing relationships between authentication attempts, and commands executed after the successful login. They collect data, synthesize it, and arrive at a judgment by weighing those factors. To reach ground truth, she asks the user, "We noticed a login from an IP address you don't normally use, followed by some unusual commands. Was that you?"

In ML, this weighing of factors is what the algorithm learns to do, but not via the method the analyst uses. Labeled examples are provided (e.g., legitimate logins, malicious logins) and the algorithm trains itself to predict which is which for cases it has not seen before. It transforms input data into numerical relationships, predicts the answer, checks whether it was right, and adjusts its calculations accordingly. Given enough time, structure, and appropriate neural network complexity for the problem, the model gets good at predicting on data it has seen before. The goal is that it then predicts well on data it has never encountered.

This is the essential point: The model produces predictions based on what it has been trained on. Understanding that is not optional for security leaders deploying these tools. It is the difference between using AI/ML with appropriate confidence and using it in ways that will eventually produce a serious failure.

Current State

Security operations today sit at the intersection of overwhelming data volume and persistent uncertainty. Even as organizations invest heavily in tooling and talent, the fundamental challenge of distinguishing authorized from unauthorized activity remains unresolved.

Security Operations: Highly Dimensional

Cybersecurity protects an organization from impact to its information and information systems by proactively strengthening the security posture of those systems. At the same time, it detects and intervenes when those systems operate in an undesirable way. While this is a concise definition, operational reality is not as clean.

The core challenge is failing to detect malicious activity. The cyber mission is making a reliable distinction between authorized and unauthorized operation across massive data volumes, in real time, with incomplete information and shifting business context. Some things are easy to observe and obviously problematic. Some are difficult to detect but unambiguously malicious once found. And some occupy genuinely ambiguous territory—visible enough, but impossible to classify without business context the analyst may not have access to. Cybersecurity has not fully solved this problem, and no amount of AI/ML adoption changes that underlying reality. This is the problem of high dimensionality. Any given security question requires an enormous number of inputs (e.g., technical telemetry, historical patterns, threat intelligence, organizational knowledge) before a reliable prediction is possible. Richard Sutton’s “The Bitter Lesson”¹ argues that computation-first approaches consistently outperform domain-specialized ones over the long run. For cybersecurity, that is both a warning and a directive: Organizations that rely solely on human expertise without scaling their computational capacity will fall behind because the data volumes involved have outpaced what human analysis can sustain unaided.

¹ Richard Sutton, “The Bitter Lesson,” March 2019, www.incompleteideas.net/IncIdeas/BitterLesson.html

Where AI/ML Helps and Where It Misleads

This brings us to the framework that should govern every AI/ML adoption decision in security operations. It comes down to how observable the activity is and how clear the authorization question is:

- **Easy to observe, clear authorization**—AI/ML excels. Signature-based detection, known IoC matching, and high-confidence automated response are appropriate. This is where automation should run with minimal human oversight.
- **Hard to observe, clear authorization**—AI/ML helps significantly with anomaly detection, behavioral baselining, and log correlation at scale. The authorization question is answerable once the activity is surfaced. Human review remains valuable, but AI/ML meaningfully reduces the detection burden.
- **Easy to observe, ambiguous authorization**—Human judgment is required. AI/ML risks false confidence here. The activity is visible, so a model trained on surface patterns may classify it with high confidence, but the classification depends on business context the model does not have. This is where over-reliance on automation produces the most dangerous failures.
- **Hard to observe, ambiguous authorization**—Neither AI/ML nor human analysis alone is sufficient. This quadrant requires integration of business context, threat intelligence, and technical telemetry that current tools support only partially. It demands the most skilled analysts working with the best available data, and it is where the gap between organizational expectation and operational reality is widest.

Observability and authorization clarity are the two variables that determine where AI/ML can be trusted and where it cannot (see Figure 2). Most SOC failures occur in the right column, where authorization is ambiguous because AI/ML models trained on technical telemetry cannot access the business context required to classify activity correctly.

	CLEAR AUTHORIZATION	AMBIGUOUS AUTHORIZATION
EASY TO OBSERVE	AI/ML Excels • Signature detection, IoC matching, automated response • Run with minimal human oversight	Human Judgment Required • Model may classify with high confidence but lacks business context • Over-automation produces dangerous failures
HARD TO OBSERVE	AI/ML Helps Significantly • Anomaly detection, behavioral baselining, log correlation at scale • Human review adds value	Neither Alone Is Sufficient • Requires integrated business context, threat intel, and technical telemetry • Demands skilled analysts

Figure 2. The AI/ML Decision Matrix for Security Operations

Consider the account creation example where a system administrator creates five new accounts on financial systems without a change control record. It clearly seems suspicious, until you learn it was directed by the CFO for a confidential external audit ahead of a venture capital round. That scenario lives in the lower-right quadrant. No model trained on technical telemetry will classify it correctly because the determining factor exists nowhere in the data the model can access. The only path to correct classification is an integrated security operation with enough organizational relationship and communication infrastructure to surface that context when it matters.

This framework is the organizing logic for the remainder of this paper.

Authorization ambiguity is where AI/ML produces the most dangerous failures. Although the activity is visible, the context that determines legitimacy is not.

Current State of AI/ML in the SOC

Security operations have always been technology focused. Ask most analysts and they will have a strong opinion about their preferred SIEM, EDR, SOAR, or case management tool. The field is generally technically proficient and curious. So why did so many SOC's find themselves caught off guard by ChatGPT's public debut and the subsequent wave of AI/GPT/LLM releases?

Machine learning is not new, PyTorch launched officially in 2017. Tensor coprocessor hardware appeared in consumer devices by 2021. The technology was available, but what was not in place was an organizational capacity to envision and mandate an integrated AI strategy for the security function.

SOCs are far more likely to select readily available tools that solve specific existing problems than to architect an integrated technology stack and enforce adoption. More capable and mature SOC's leverage integration with emerging IT development projects as the funding vehicle for technology exploration. Most SOC's do not have this IT integration established at all, let alone refined and optimized.

What changed with ChatGPT was not the technology, it was the popular imagination of what was possible. Vendors responded with an abundance of AI-enabled and AI-native products. Security operations teams began using them, with or without organizational sanction. This disruption and resulting confusion is readily apparent everywhere now.

Current AI Technology Use

We know from the 2025 SANS SOC Survey² that employees choose a technology to accomplish work for their convenience, not because the organization has intentionally selected and deployed the tool in an intentionally secured way. AI tools are not part of the defined workflow, but staff are using them anyway. That is the definition of shadow IT, technology adopted for individual convenience rather than through intentional organizational selection and deployment. The security implications of that pattern in a security operations context specifically doesn't need an explanation.

Vendor-provided AI integrations are arriving through a different channel: software updates. These capabilities are showing up across the core tool set security teams already rely on:

- SOAR tools are receiving enhanced decision-making capability for automations and orchestration.
- Knowledge management platforms are gaining AI-powered chatbots.
- Ticketing systems are implementing AI-assisted assignment and closure. Research consistently shows that initial ticket assignment to the right person is the strongest predictor of successful resolution within expected service levels.³

AI-assisted assignment is a legitimate and relatively low-risk application of the technology in this context.

Current Staff Skills and the MCP Gap

Detection engineering represents the highest-potential application of ML in the SOC, and the area where most teams currently lack capability. Most SOC's do not have the MLOps expertise to engineer their own detections, and the availability of GPT-assisted programming may change that. A mid-career or senior analyst with access to a dedicated workstation could prototype internally trained ML detections using purely internal data, combining the broad-visibility detections that larger vendors provide with custom algorithms tuned to the organization's specific environment.

Realizing that potential requires staff who can develop and implement AI tooling and move beyond consuming vendor products. Most vendor integrations plug directly into existing deployments, but enabling those products to leverage retrieval-augmented generation (RAG) or agentic operations requires working with the model context protocol (MCP), the open source specification released by Anthropic in 2024⁴ that has rapidly become the de facto standard for enabling models to access external data and act on it.

MCP introduces real security risk. When the company that builds the underlying AI tool can accidentally expose its own developer product's source code (as Anthropic did with Claude Code) the probability that security staff will make similar mistakes with organizational data is meaningful. This is not a reason to avoid the technology but instead a reason to approach it with the same rigor applied to any other security-critical infrastructure.

² SANS 2025 SOC Survey, www.sans.org/white-papers/sans-2025-soc-survey

³ "Intelligent Ticket Assignment System: Leveraging Deep Machine Learning for Enhanced Customer Support," April 2025, www.researchgate.net/publication/391847347

⁴ "Introducing the Model Context Protocol," November 2024, www.anthropic.com/news/model-context-protocol

Immediate Value of GPT

The highest-return near-term application of GPT technology in security operations is not detection but rather communication. Security staff are consistently poor at quantifying and articulating the value of security operations and the business impact of undesired activity. A GPT configured with a defined organizational schema for impact evaluation can dramatically improve the quality of reporting to internal stakeholders and the contextual information shared with MSSP partners.

That reporting does not need to flow only from SOC to business. A well-designed GPT integration could automatically provide the MSSP with contextual insight about asset stakeholders, potential business impact, and organizational risk appetite, improving the quality of MSSP investigation and the relevance of what comes back. The static ticket model gives way to a collaborative, iterative exchange between systems and analysts that gets the right information to the right people in the right format.

Near-Term Changes

Most security operations remain reactive by necessity, responding to a threat landscape shaped by frontier AI vendors who prioritize rapid capability development over security stability. The near-term path forward is to build an operational model that holds rigorous engineering and rapid preliminary assessment together, in sequence rather than in tension.

Proposed Future States

Almost every security operation is in reactive mode. The frontier model vendors (OpenAI, Anthropic, Google DeepMind, xAI, Mega, Mistral AI, DeepSeek, and others) are in a period of rapid, competitive feature development that prioritizes capability over security.

We've been seeing this rapid technology deployment for a while. IPv6 deployed a new network protocol standard that no one was ready for. DNS over HTTPS substantially reduced network monitoring visibility because browsers rolled it out unilaterally. The appropriate response is not to wait for stability that will likely not arrive. It is to build an operational paradigm that holds both things at once: the rigor of thorough engineering and the agility to make rapid preliminary assessments when time does not allow for more. Those are not in conflict. Instead, they are sequential: assess quickly, then continue the longer engineering effort in parallel.

In a world where “move fast and break things” is a popular slogan, “be smooth and fix things” is the cybersecurity response. Not because caution is more comfortable, but because smooth and corrective are what “move fast and break things” wants to be. It wants to ship, iterate, and improve. Security operations do that too, with the added constraint that some failures are not recoverable. The operational discipline involves constant reassessment and enhancement, not deployment and observation.

In the near term, there is no excuse for failing to deliver clear articulation of security posture to all relevant audiences. The capability to do so exists today.

Security Infrastructure as Code

As teams transition to a paradigm where AI tools require MCP and other data transport between models, they are entering the realm of security infrastructure as code. The failure modes are specific and serious. Inappropriate handling of security data produces loss of confidentiality, integrity, or availability, the very outcomes the security operation exists to prevent. Indiscriminate data ingestion and transport produces unnecessary cost without improving detection capability.

Many SOC's are stuffing data into their SIEM because they have not done the engineering work to understand what data their detections require. The discipline of data selection—collecting and transporting only what is needed for specific detection and analysis purposes—is an engineering problem. It is also the foundation on which effective AI-enabled tooling must be built. A model is only as good as the data it can access, and indiscriminate data collection does not compensate for poor detection engineering.

The Hunting-to-Detection Engineering Pipeline

The operational model that best positions a SOC for AI integration is built around a clear pipeline: Hypothesis-driven hunting feeds detection engineering, which feeds automated detection and response (see Figure 3).

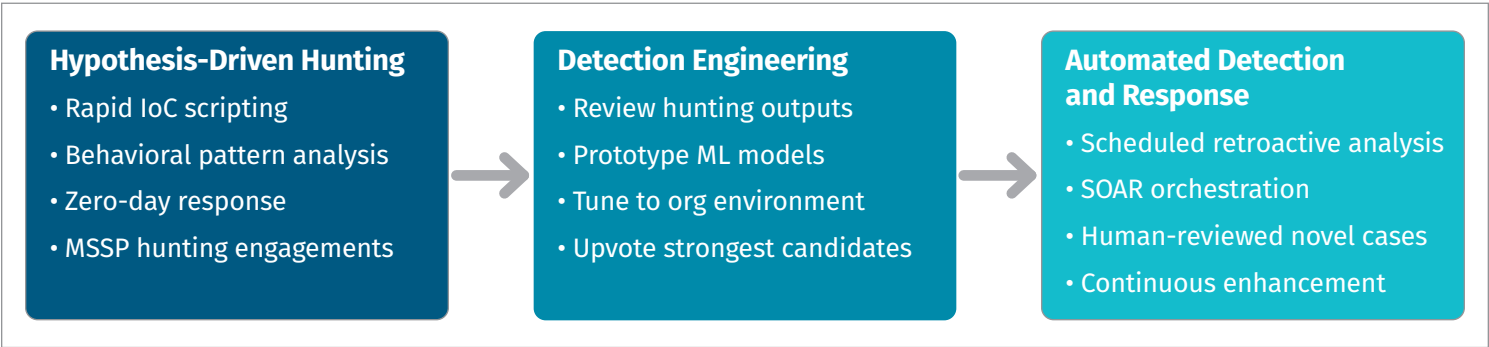


Figure 3. Hunting-to-Detection Engineering Pipeline

Hunting should be rapidly adaptive, flexible, and inexpensive to execute. For example, when a limited-scope zero-day vulnerability is disclosed with few public indicators, a GPT can be used to rapidly draft a PowerShell script to search Windows event logs for available IoCs. The first draft will require iteration and, once usable, the script runs across all potentially affected systems. When additional indicators become available, the script updates and reruns. This kind of retroactive IoC analysis can be fully automated and run on a regular schedule.

More sophisticated hunting goes beyond known IoCs to hypothesis-driven analysis of living-off-the-land techniques and behavioral patterns that do not yet have defined indicators. This is the work that requires experienced analysts, and it is the work that should feed the engineering pipeline. Hunts that surface promising detection opportunities get upvoted by hunting staff and prioritized by detection engineering. The engineering investment is reserved for the strongest candidates from the hunting pool.

Teams that cannot execute this internally can negotiate structured hunting engagements with their MSSP, typically once a month if the contract is written appropriately. Larger MSSPs maintain high-powered analysis workstations capable of rapidly prototyping custom ML models to review data at scale.

The analyst-developer-engineer is not a future role. It is the role the near-term SOC needs now.

The change that needs to be made is to integrate the MSSP as an extension of internal capabilities. This can be done without violating data compartmentalization and control. But it entails sharing context, enrichment, analytical hypotheses, rationales for decisions, and root cause assessments. The opposite of this is sending a ticket or alert with little to no information, then closing the ticket without explanation. This context-deprived communication swells frustration and eliminates trust.

Effective execution of this pipeline requires staff with a specific combination of expertise with IT systems operations and architecture, vulnerability management, detection analysis and engineering, incident handling, and threat intelligence. It also requires conceptual fluency in DevSecOps and MLOps, including deployment options such as APIs for detection, MCP, and RAG. The analyst-developer-engineer is not a future role; it is the role the near-term SOC needs now.

Analysis Rigor and the Limits of Automation

A persistent deficiency in security operations is analytical rigor, the discipline of structured, repeatable analysis that minimizes individual bias and improves consistency. Security operations should not attempt to apply the full scientific method, but rather a simplified analytic technique modeled on Analysis of Competing Hypotheses: Enumerate hypotheses, identify evidence that confirms or disconfirms each, and assess the weight of evidence before concluding. This approach reduces bias, improves repeatability, and maps well to workflow automation and orchestration.

Automation is the programmatic repetition of a task and orchestration is that automation across multiple systems. The concept predates SOAR tooling. Experienced administrators called it scripting. SOAR provided scaffolding and normalization. AI-enhanced SOAR promises objective performance improvement without requiring every step to be programmatically defined.

The risk is that AI-enabled automations will make mistakes that are difficult to anticipate, startling in their consequences, technically logical, and operationally disastrous. An automation that wipes a hard drive to eradicate a confirmed compromise has a certain perverse logic. It is not a decision any security team would sanction if asked. Maintaining human review of automated action in novel situations is not timidity. It is the appropriate response to the paradox of automation. As systems become more complex and automated, the humans responsible for intervening when things go wrong become less capable of doing so, precisely because of that complexity.

Building human judgment into the automation pipeline is how the SOC hedges against that risk. Leveraging the MSSP to provide industry-wide visibility and an objective critical eye is an obvious enhancement and a valuable use of an available resource. Operationally implementing this is incumbent upon every organization.

Reporting and Communication

In the near term, there is no excuse for failing to deliver clear articulation of security posture and issues to all relevant audiences. Well-designed metrics, effective visualizations, and audience-appropriate reporting are achievable now. A GPT configured with organizational context can produce real-time queries, visualizations, and scenario modeling for any authorized user. The MSSP alert that currently arrives as a ticket can become the start of a collaborative investigation conducted through a chat interface, querying both the MSSP's limited external view and the organization's richer internal data simultaneously, without those sources appearing as separate disconnected inputs. That capability is available today for organizations that choose to build it.

Long-Term Changes

The long view on cybersecurity is about structure, not which threats emerge or which tools prevail. The organizational separation that allowed security to mature as a discipline has run its course. What comes next is integration, with security expertise embedded across operations, development, and business functions rather than housed in a parallel structure alongside them. AI/ML tools do not cause that shift, but they make it viable at the scale modern organizations require.

Pole-Star: Your Long-Term Navigation Aid

The longer-horizon prediction for cybersecurity is convergence. The specialization of security operations into a distinct organizational function was a necessary phase because cybersecurity needed time and organizational separation to mature as a discipline. That maturation has occurred and the way forward is integration.

In the early years of this century, cybersecurity became a necessity for IT operators. It became a specialization because availability consistently won the argument against confidentiality and integrity in operational trade-offs. The resulting separation allowed security practices to develop without being constantly subordinated to uptime pressures. But that separation has produced its own costs: the MSSP visibility gap, the organizational distance between security teams and business context, and the persistent failure to integrate security into development and operations from the start.

The parallel is instructive. Networking was once a deep specialization. So was virtualization. Both reintegrated into general IT operations as tooling matured and the knowledge required became widely distributed. Cybersecurity is on the same trajectory. The AI/ML tools now emerging accelerate that convergence—not by making security easier, but by making it possible to operate at the scale and speed that full integration requires.

This does not mean security becomes less important, but it does mean more work, done by staff who hold both security and operational expertise simultaneously. The paradox of automation makes this more urgent. As automated systems become more capable and complex, the humans responsible for them need both deeper and broader knowledge. The answer to increasing automation is knowledgeable people with better tools, visibility, and time to do appropriate analysis.

Staff Investment

Your staff will make the difference between being able to customize and develop cybersecurity tools for your own organization or accepting the as-offered tools from vendors. Simpler automations and customizations get done internally or under contract to the MSSP, preserving information trust boundaries while ensuring that organizational intent, not vendor default, governs how the security infrastructure behaves. More sophisticated development decisions get made by staff who understand both the technical options and the organizational risk appetite well enough to choose correctly.

Making that argument to executive or board-level decision makers requires a clear argument around self-directed agency over information systems being smart business. Organizations that accept vendor defaults wholesale cede control over systems that carry their most sensitive information. That is not a security argument alone, it is a business continuity argument, a competitive argument, and an organizational resilience argument.

The security professional of this future operates as an amalgamated analyst-developer-engineer, working within an AI-assisted development environment.

Selecting an MSSP

To start with the basics, the key things to look for in an MSSP include alignment of capability, technology, information provenance, and workflow compatibility.

As an example of capability, the specific industry vertical should already be addressed by the MSSP. If they're not already working in your vertical, patience and training are required by the organization. This is reasonable when the organization operates in a vertical rarely outsourced, but it sees the long-term value of developing the capability in the MSSP.

Having the exact same technology as the MSSP isn't a requirement, but it makes implementation substantially easier. In cases where the MSSP doesn't already support a specific technology, ensure in the contracting that it will be absorbed or excluded as appropriate.

Information provenance is abiding by legal requirements. Some locations or industries mandate that data stay within a jurisdiction or require specific controls for information leaving a jurisdiction. Align this early in your MSSP selection process.

Finally, think about how your team currently works and how it would like to work. Look for an MSSP that fits that model or be willing to adjust your operations to conform to the MSSP's best practices.

The MSSP Relationship: What It Can and Cannot Do

Organizations overwhelmed by the difficulty of distinguishing authorized from unauthorized activity at scale frequently turn to MSSPs or managed detection and response providers. These entities offer real value in breadth of visibility across large customer bases, economies of scale for hunting and detection engineering, and staff with deep analytical expertise. They excel at finding things that are easy to observe and clearly malicious.

The prospect of transferring work to an already knowledgeable specialized organization at a lower cost than could be done by the organization has substantial appeal. The prospect of leveraging that partner to assist with the AI problem is an obvious add-on to the existing efforts already transferred to an MSSP.

The MSSP's limitations stem from a structural constraint in that the MSSP has less visibility than the organization itself. Organizations limit the data they transfer to an MSSP because every bit costs money to send and store. That constrained view makes nuanced discovery harder. And for the authorization ambiguity problem of determining whether a given action is legitimate given the organization's specific business context, the MSSP is structurally unable to help. Only the organization can authoritatively answer that question.

Using an MSSP is not purchasing an insurance policy. It is leveraging the resources of a closely collaborating partner to investigate and assess sensitive actions across all organizational IT systems. The customer must continue to evaluate authorization, operate at the tempo the MSSP has established, and return insight as a collaborative effort. Organizations that approach the relationship otherwise will be disappointed—and will likely blame the MSSP for limitations that were always structural.

A good MSSP will enable you to have greater confidence in analysis by producing defensible decisions quickly, scaling to the size of your organization and information assets. This is the same objective we currently have for AI/ML, and the MSSP can assist with that effort.

Conclusion

The current state of AI in security operations is a rapidly shifting collection of tools being used in undisciplined ways, producing reasonable but inconsistent results. The near-term path forward re-emphasizes the value of knowledgeable staff, builds the hunting-to-detection engineering pipeline that positions teams for genuine AI integration, and establishes the reporting and communication infrastructure that makes security operations legible to the rest of the organization. The longer-term vision is a security function that has fully converged with IT operations, not as a diminishment of security, but as its maturation.

None of this happens without intentionality. The pressure to adopt AI tools is real and will not abate. The risk of adopting them without a framework, without understanding what they can and cannot do, without preserving human judgment where the authorization question requires it, without building the staff capability to operate and adapt the tools over time, is equally real. Massive churn with little positive change is the likely outcome for organizations that respond to pressure without direction.

We are at an inflection point. The decisions security leaders make now about tools, staff, and the relationship between AI capability and human judgment will define their organizations' security posture for the next decade. The power to make those decisions well with intentionality, vision, and a genuine understanding of what the tools do, is available. This paper is intended to help with that. The key points are to:

- Understand AI and that it is churning and changing currently.
- Select and engage MSSPs effectively for high-value operations.
- Recognize that AI helps visibility, observability, and authorization; but not in the same way for each organization.
- Remember that AI tools enhance communication, automation, analysis, and detection engineering.
- Look for incremental enhancements.
- Be flexible with your AI deployments.
- Staff with cyber-first principles—AI for cyber, engineering, and MSSP/outsourced engagement are scarce, so find and retain them.

Sponsor

SANS would like to thank this paper's sponsor:

TATA COMMUNICATIONS