

From AI-powered tools to an AI-powered SOC

How you can turn people, process, and technology into
an innovative operating model for frontier firm SOC's



What's inside

1. *The world is changing and the balance has shifted*
2. *Capgemini's approach to new SOC processes: AI vs automation vs human analysis: three essential selection criteria*
3. *Four steps to deliver successful AI SOC deployments*
4. *How to avoid common AI SOC pitfalls*
5. *The AI-powered SOC in practice*
6. *The AI-powered SOC: A Capgemini solution, leveraging Microsoft Security*
7. *Appendix: AI use case assessment prompt*

As AI accelerates cyber threats, traditional approaches built on more tools, more alerts, and more analysts are reaching their limits.

This eBook explores how organizations can transform security operations through a governed operating model that combines AI, automation, and human expertise to improve resilience, efficiency, and response outcomes.



The world is changing and the balance has shifted.

cyberattackers now have the ability to capitalize on exploits faster than ever before.

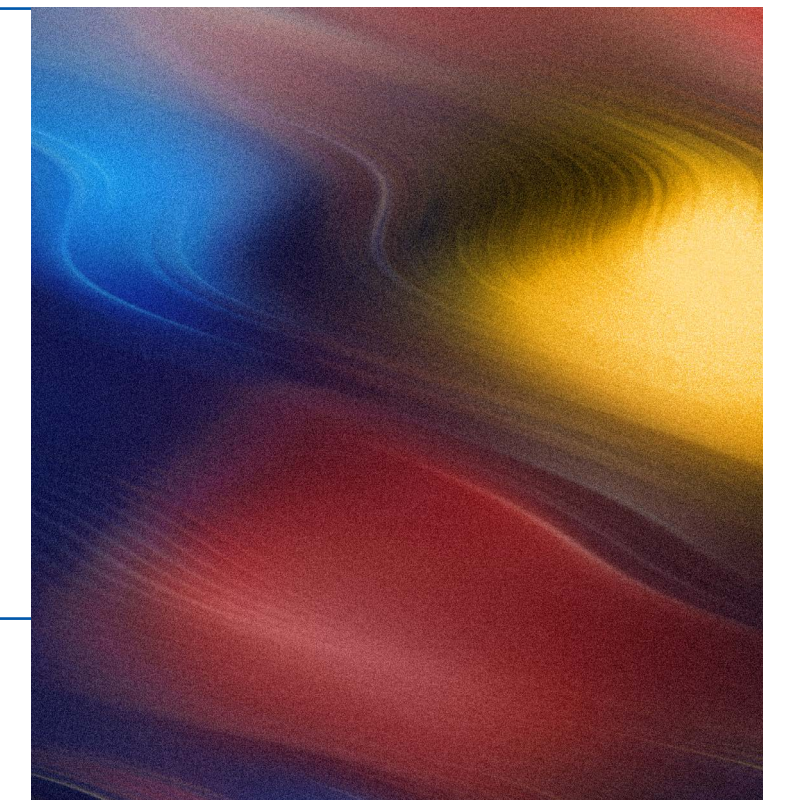
In 2022, the average time to exploit was **9.7 months¹**. By 2026, a few years after the public release of Gen AI tools, that number had fallen to just **nine hours²**. Estimates suggest AI agents will further reduce the time to exploit by **50% by 2027³**.

The question remains: how do you stop AI-powered cybercriminals? For many, the answer is to deploy AI in the SOC. However, not every organization is adopting AI security tools at the same rate. Some SOC's are ready to deploy, while many more are taking a wait-and-see approach to avoid potential pitfalls and security risks.

It's easy to see why there's hesitation, as poorly managed AI use cases can create more problems than they solve, and add to the cost, complexity, and risk of cybersecurity. However, by selecting the right AI use cases and implementing them carefully, SOC's can help security analysts do more with less, and detect growing threats faster.

In 2025, Microsoft Detection and Response Team researchers uncovered an unusual backdoor. The exploit, called **SesameOp**, involves using the OpenAI Assistants API as a C2 channel to covertly communicate and orchestrate malicious activities within a compromised environment.

This is just one example of the new attack vectors and cybersecurity risks opened up by LLMs and generative AI platforms. It's far from the only threat though. Armed with these AI tools, all



¹Zero Day Clock, From Vulnerability to Exploitation Dashboard.

²Zero Day Clock, From Vulnerability to Exploitation Dashboard.

³Gartner Predicts 2025: Navigating Imminent AI Turbulence for Cybersecurity.

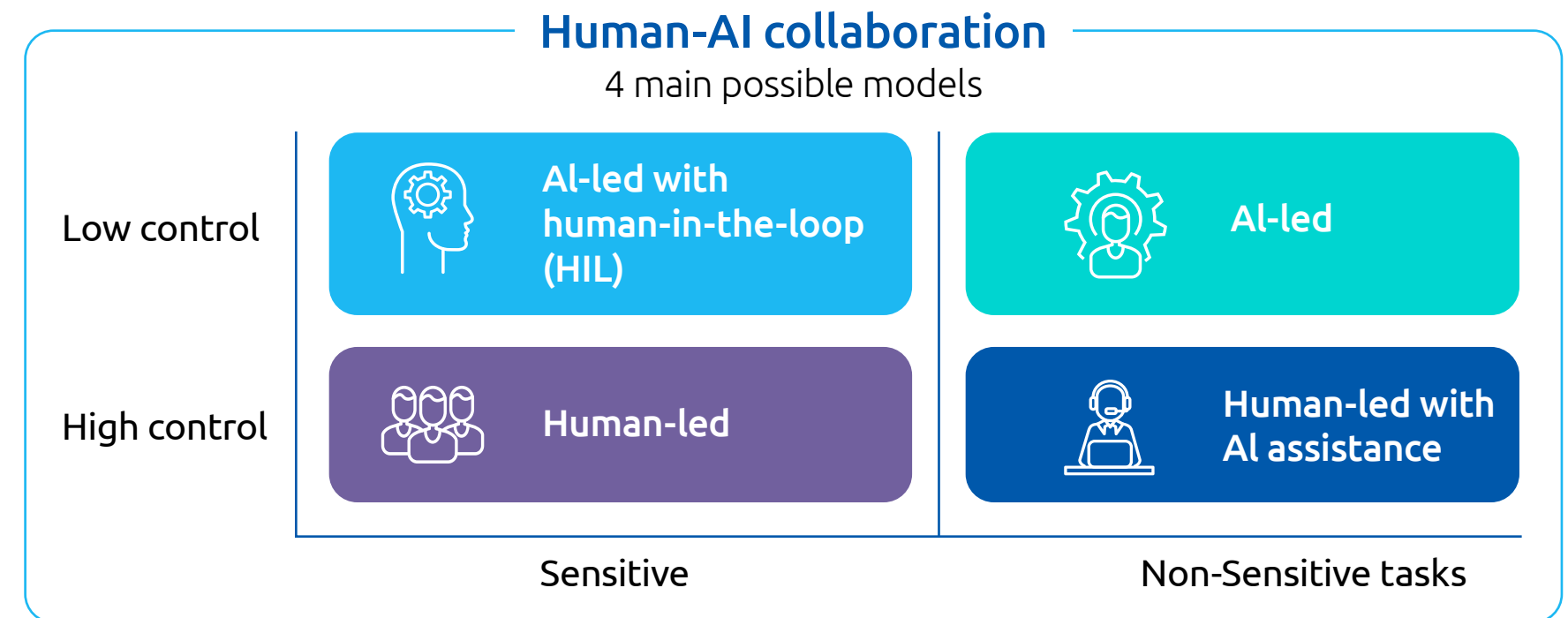


Capgemini's approach to new SOC processes: AI vs automation vs human analysis: three essential selection criteria

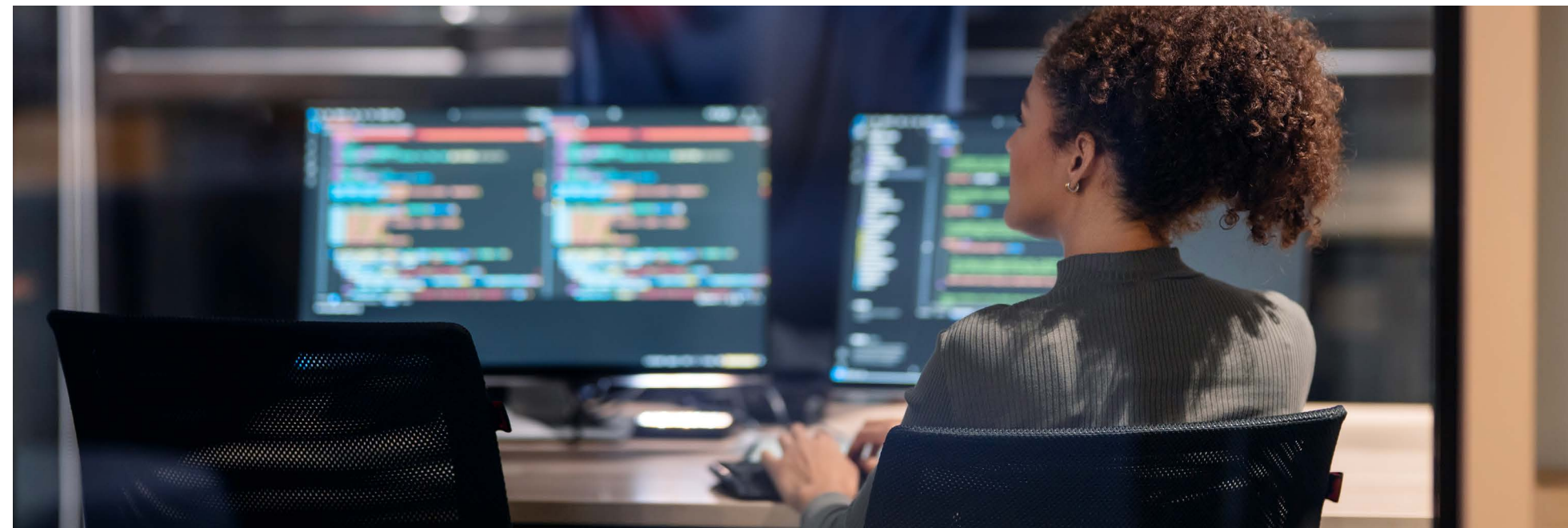
There are three essential selection criteria for AI use cases in the SOC: the collaboration model, the compute cost and frequency, and resource management planning.

Choosing the collaboration model

AI tools are powerful engines for synthesizing and summarizing unstructured data, but they can come with drawbacks depending on implementation. AI and human collaboration exists across four quadrants.



Depending on control and sensitivity of the use case, AI will be less effective or more expensive than turning to human analysts or more traditional automation.





To identify the best approach for a given collaboration use case, consider these four criteria:

- **Determinism:** Does the task have a clear, repeatable logic that would be better served by traditional automation?
- **Judgment:** Does the use case involve synthesizing, summarizing, or sorting large quantities of unstructured data?
- **Auditability:** Can you ensure AI prompts and outputs are traceable and repeatable?
- **Frequency:** How regularly will AI functions need to be called to complete the task, and how much will that cost?

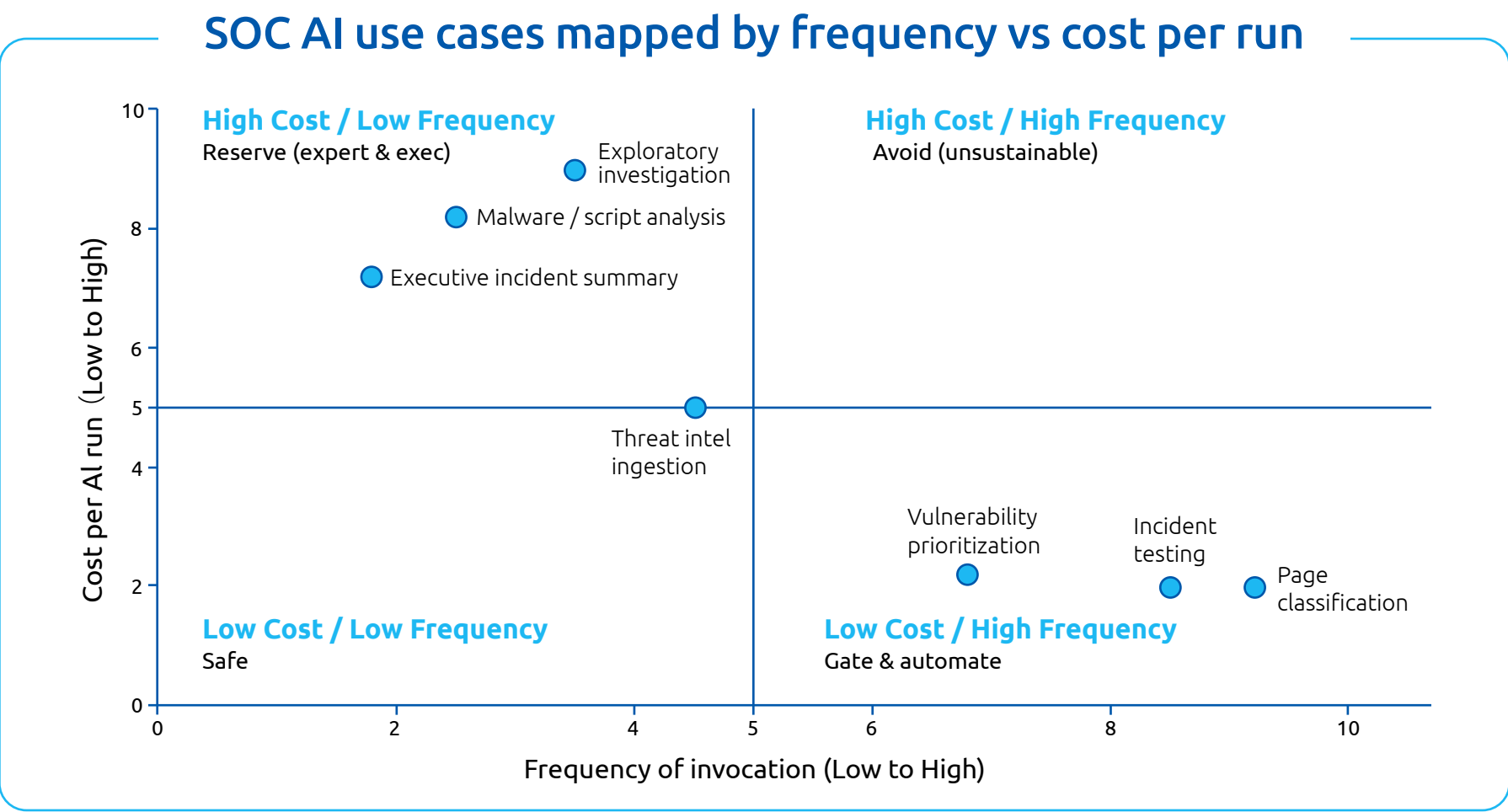
Depending on the balance of these factors, it becomes much easier to measure whether a task is best automated, covered by AI but with strong gates and controls in place, or should be completely engineered out of your process.

See Appendix A for a starter AI prompt that can assess use cases against these criteria to speed up your process.

Assessing compute costs and frequency

Compute Costs are a crucial factor, as frequency and raw credit demands both affect how quickly a use case will drain limited AI resources.

For example, high-cost, low-frequency tasks such as malware analysis and executive incident summaries are appropriate for AI, but are best reserved for scenarios involving significant human intervention. Meanwhile, high-frequency, low-cost tasks such as phishing classification and vulnerability prioritization can be automated, but with gated AI segments used where appropriate.



Resource management planning

Unstructured Microsoft Security Copilot usage is expensive because AI costs scale with prompt length, chaining, and repetition, as mentioned above. Analysts tend to “chat” for reassurance. Each variation consumes fresh compute.

In the case of Microsoft Security Copilot, all Microsoft 365 E5 and E7 customers get access to 400 Security Compute Units (SCUs) per month for every 1,000 paid user licenses (**up to 10,000 SCUs a month**). Some things are worth it if they’re free, some are not. Picking up on your firm’s SCU patterns and usage offers a key indicator of what use cases to prioritize.

Different AI use cases will go through credits at different rates:

 Low-cost AI tasks include single-prompt tasks	 Medium-cost tasks They often require of 3–5 prompts chained together with no automation	 High-cost workflows such as malware analysis and exploratory investigations involve multi-step reasoning that demands more AI capacity.
--	--	--

Workload multipliers also add or change these costs—examples include things like:

- | | |
|--|---|
| <ul style="list-style-type: none">• Large log volumes• Script or malware analysis• Executive reporting and narrative personas• Prompt syntax complexity | <ul style="list-style-type: none">• Context schemas• Guardrails and constraints• Confidence thresholds• Human-in-the-loop rules• Expected outputs |
|--|---|

However, you can bring costs down by pre-querying data in Kusto Query Language, spinning resources up and down, pre-loading guardrails and constraints outside prompts, establishing thresholding (e.g. don’t run if within x% of capacity), and by caching and reusing outputs to scale.



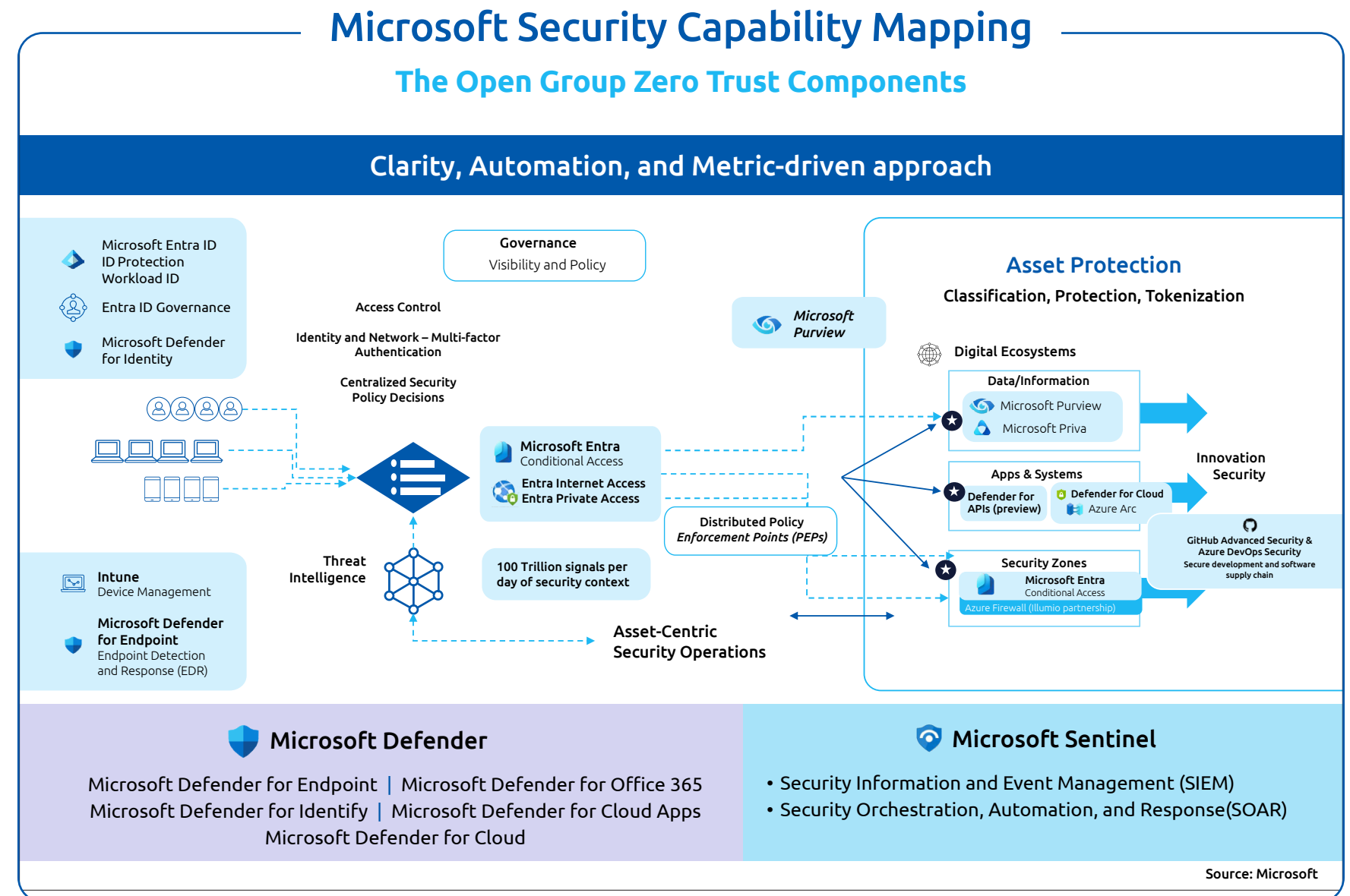
Four steps to deliver successful AI SOC deployments



Selecting the right AI SOC use cases sets the stage, but it doesn't guarantee success. Follow these four steps to create a smooth transition to a secure, efficient AI-powered SOC.

Step 1: Establish a Zero Trust baseline

AI doesn't improve the way your SOC operates by default. But it will help disciplined SOC's scale up. That makes it essential to establish strong principles before you introduce any AI. And strong cybersecurity principles begin with a foundation of Zero Trust.



Tools like **Microsoft Zero Trust Assessment** can automate the process and check your security configurations against Secure Future Initiative and **Zero Trust pillars**. And the right cybersecurity partner can help you turn this initial assessment into a comprehensive security roadmap.



Step 2: Map your value streams

Inefficient processes in your SOC run the risk of bloating and eating away at your AI resources if they aren't addressed before establishing AI use cases. Value stream methodologies such as **Six Sigma** can help identify these inefficiencies so you can engineer them out.

Dimension	Why it matters
Actor (human/system)	Ownership clarity
Tool(s) used	Consolidation targets
Touch time	Labor cost driver
Wait time	MTTR driver
Handoffs	Error and delay risk
Automation rate	AI readiness signal

This is a crucial step in controlling costs, as any automation applied to an inefficient or unnecessary process is resource that could be better spent elsewhere. However, the roles you have in the SOC will not fully align with what you'll need for the SOC of the future.

While some roles will remain similar, industry researchers like Gartner's Chief of Research for Cybersecurity, Pete Shoard, also **expect new roles will join the SOC**. This includes roles like expanding the scope of detection engineers specialize in prompt design and detection development around AI and offensive testers that will be needed to code offensive exercises or validate exposures will join the SOC ranks. Many roles will also shift to validating secure code infrastructures for AI rather than manually analyzing incidents. This change also has downstream implications for training, career growth, and staffing levels.

Roles of a SOC Team in 2030

SecOps Manager			• Metrics/reporting • Communications • Operations tasking • Coordination		
 Detection engineer • Prompt design • Threat/technique assessment • Detection development • Operational feedback management • Data validity • Workflow design/test	 Systems engineer • Data parsing • API integration • System efficacy/efficiency • Alert/case accuracy/efficacy • Data validity	 Threat expert • Threat intel evaluation • Threat hunting • Offensive scenario design	 Threat analyst • Incident management • Response coordination • Executive communication • Lessons identified analysis	 Exposure specialist • CTEM program execution • Mobilization coordination	 Offensive tester • Red-team exercise design/execution • Penetration testing • Offensive exercise codification

Zero Trust in practice: identify entitlement management

Identity is a key Zero Trust pillar. Entitlement management is the discipline of defining, granting, governing, and removing what a user, service, agent or partner is allowed to access and do across systems, applications, and data.

AI workloads are powered by non human identities, such as agents, service principals, or applications, yet they're often granted broad, static permissions and don't have the same entitlement processes humans have. The result is silent over entitlement, weak auditability, and a large blast radius if something goes wrong.

Traditional role-based access control and manual reviews can't keep up with the speed of AI adoption when queues are flooded with access reviews from every agent. Microsoft explicitly positions AI agents as first-class identities that shouldn't have ad hoc exceptions. Instead, they require the same lifecycle governance as humans. They also require a context layer of intent to validate access.

The only way to govern access is to use risk and intent signals from Microsoft to enable role-based, time-bound, and continuously reviewed access and authorization process. Microsoft Entra features enable programs to achieve:

- **Microsoft Entra Entitlement Management** automates access using access packages, approval workflows, time-bound access, and expiration.
- **Microsoft Entra Agent ID** inventories agents, tracks what they can access, and enables least-privilege controls to prevent agent sprawl.
- **Microsoft Entra Privileged Identity Management (PIM)** enforces just-in-time privileged access for both humans and agents, eliminating permanent high-risk permissions.
- **Access Reviews** continuously verify whether access is still required, creating audit-ready oversight for AI-driven environments.

Capgemini's IAM Assessment, aligned to Microsoft FastTrack methodology has repeatable, packaged processes that combine these features to determine where to harden some of the most important attack surfaces.



Step 3: Understand your SOC's Maturity

Almost all AI use cases are better when they're supported by rigorous standard operating procedure (SOP) and governance models. AI doesn't make a bad SOC good. It makes a disciplined SOC scale. Capgemini aligns to CMM as it excels at operationalizing the NIST Risk Management Framework, specifically for security operations.

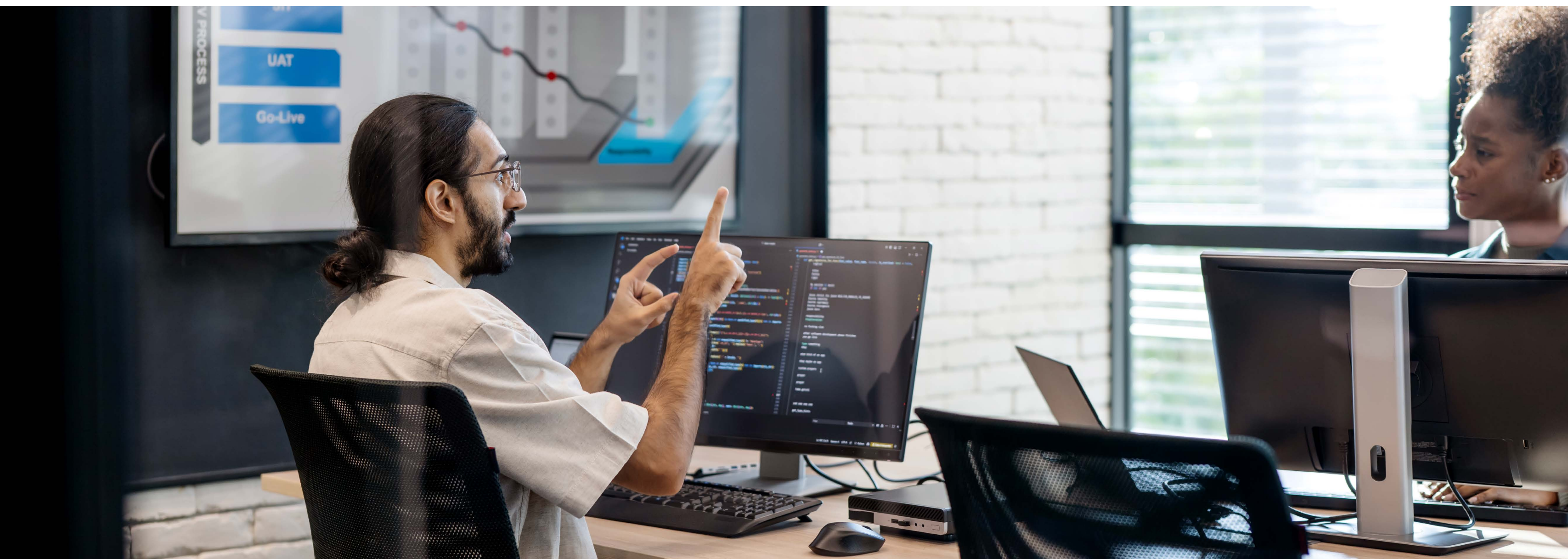
The **SOC Capability Maturity Model** (SOC-CMM) helps assess governance, SOP, and other core parts of security operations against industry baselines. By improving your SOC-CMM maturity,

you can track the improvement and establish processes and guardrails that keep your SOC operating and growing against a shared vision of excellence.

Choosing SOC-CMM allows a broader framework like NIST to crystallize even further within the SOC domain and simplifies compliance with NIST audits. For clients with multiple operations centers, a detailed framework also creates repeatability and removes ambiguity around successful and compliant SOC operational processes.

SOC CMM is the ideal way to anchor AI planning to SOC readiness by ensuring operating procedures and orchestration layers are aligned and functional.

SOC domain	What it covers
 GOV Governance	Ownership, policy, risk posture, audit readiness
 ID Asset	Assets, identities, threat intelligence
 PR Preventive Controls	Hardening, identity, vulnerability reduction
 DE Detection Engineering	Telemetry, detection logic, coverage
 RS Response & Orchestration	Triage, containment, recovery
 IM Improvement Metrics	Metrics, lessons learned, continuous improvement





Step 4: Establish the building blocks of AI functionality

In modern SOC's, detections have shifted from static rules to detections-as-code. Detection logic is written, versioned, tested, and deployed through pipelines. This gave us peer review, consistency, rollback, and measurable coverage, similar to application code. However, detections only cover one part of the SOC's workflow.

Promptbooks in an AI-powered SOC are another opportunity to bring software development lifecycle and software factory discipline directly into security operations.

Instead of ad hoc analyst judgment or opaque AI behavior, workflows are encoded as version-controlled promptbooks that act like reusable code modules. Each promptbook captures and references intent, inputs, logic, guardrails, and expected outputs, and is tested, reviewed, and updated through CI/CD-style workflows. Most importantly, they can be chained to other parts of the workflow using automation.

This approach introduces peer review, change control, rollback, and auditability to SOC decision-making in the same way modern software teams manage application code.

AI Powered SOC, detection logic, triage decisions, enrichment steps, and response actions become composable “building blocks” that can be validated, improved, and reused across use cases. The result is a SOC that evolves through iterative releases, measurable quality improvements, and controlled change, rather than one-off playbooks or experimental AI deployments.

However, before you begin creating and assembling workflows, there are key design principles to keep in mind:

- Automation before AI
- Zero Trust by default
- One decision point per severity level
- Where ambiguity exists, rely on human judgment only
- Tie metrics to outcomes, not activity
- Smaller pieces are delivered faster

For example, you might build a monolith of code, but this suffers from the “Christmas tree light problem”, when one light burns out, all of them stop working. On the other hand, you could build multiple blocks of code to separately handle:

- Underlying data logic, including Python script execution, LLM interpretation, and error handling
- SOC workflows, including variable-supported prompts and automation input
- Synthesis, to create deterministic outputs in tasks like auditable task records

Blocks like these can enable agentic capabilities for specific micro-tasks in your SOC workflow, which can then be connected to each other as reusable components.

Because each block can be configured to allow for different retry counts and intervals, identifying and reflecting failure codes and details are also included to ensure flexibility. Modular blocks also make it easier to create a prioritized backlog of additional AI-driven workflows.





How to avoid common AI SOC pitfalls

On your journey toward an AI-powered SOC, it's worth keeping an eye out for common AI integration pitfalls that can stall progress and harm security outcomes.

Common pitfalls include:

- Inconsistent use cases across users and teams
- Heavy dependence on individual skill and prompting abilities
- Limited governance and auditability
- Inconsistent prompts that increase the risk of hallucinations

Navigating the AI SOC pitfalls

We recommend three approaches to work around these pitfalls and find more consistent AI SOC success.

1. Deploy a promptbook library

Promptbook libraries provide a governed, reusable, and role-aware way to standardize how AI systems are used across your organization. Instead of ad hoc prompting, teams are given prompts as a managed asset, similar to code libraries or runbooks.

A well-designed library should ensure that:

- Promptbooks are stored as versioned JSON artifacts
- No credentials, secrets, or execution logic are stored in the repo
- The repo is the authoritative source of truth for proposed state only

Each workflow in the library follows a block schema:

- 1. Detection/trigger** (automation)
- 2. Enrichment** (automation-first)
- 3. Analysis and synthesis** (human + AI)
- 4. Disposition** (human + AI gated)
- 5. Response** (automation, post-approval)
- 6. Summarization** (AI-assisted, structured)



2. Trigger promptbooks with automation

Microsoft Security Copilot uses REST based plugins, making Microsoft Azure Foundation a natural integration point for calling Security Copilot or acting as a proxy for action triggers. This offers an ideal place to automate trigger promptbooks while prioritizing resources more effectively and offering more consistent costs.

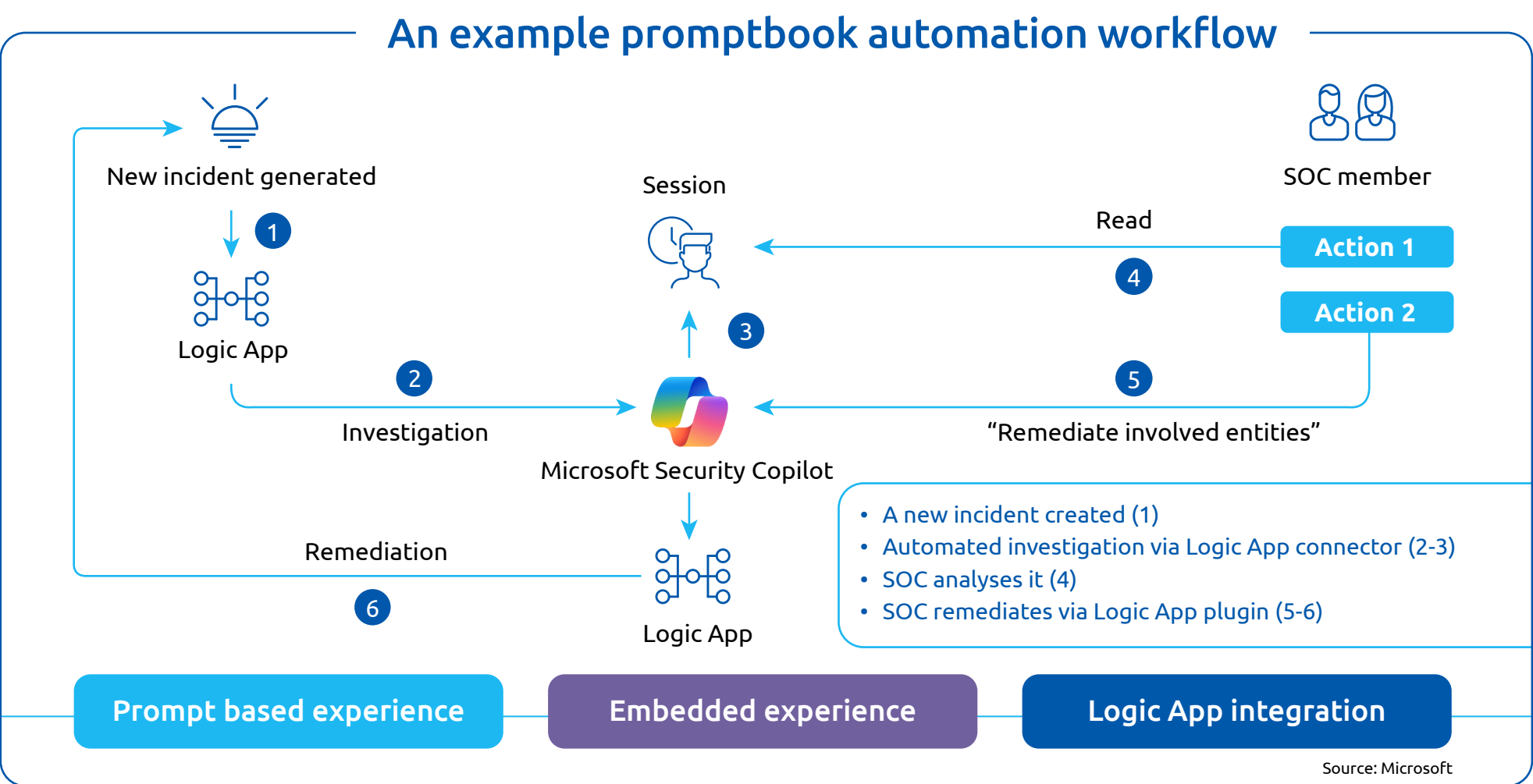
3. Design prompts for degradation, not perfection

Unless told otherwise, AI tools always assume they're working with ideal data, stable telemetry, and clean inputs. The problem is that most SOC failures don't happen during the ideal scenario.

That's why it's important to make sure your promptbooks clearly define:

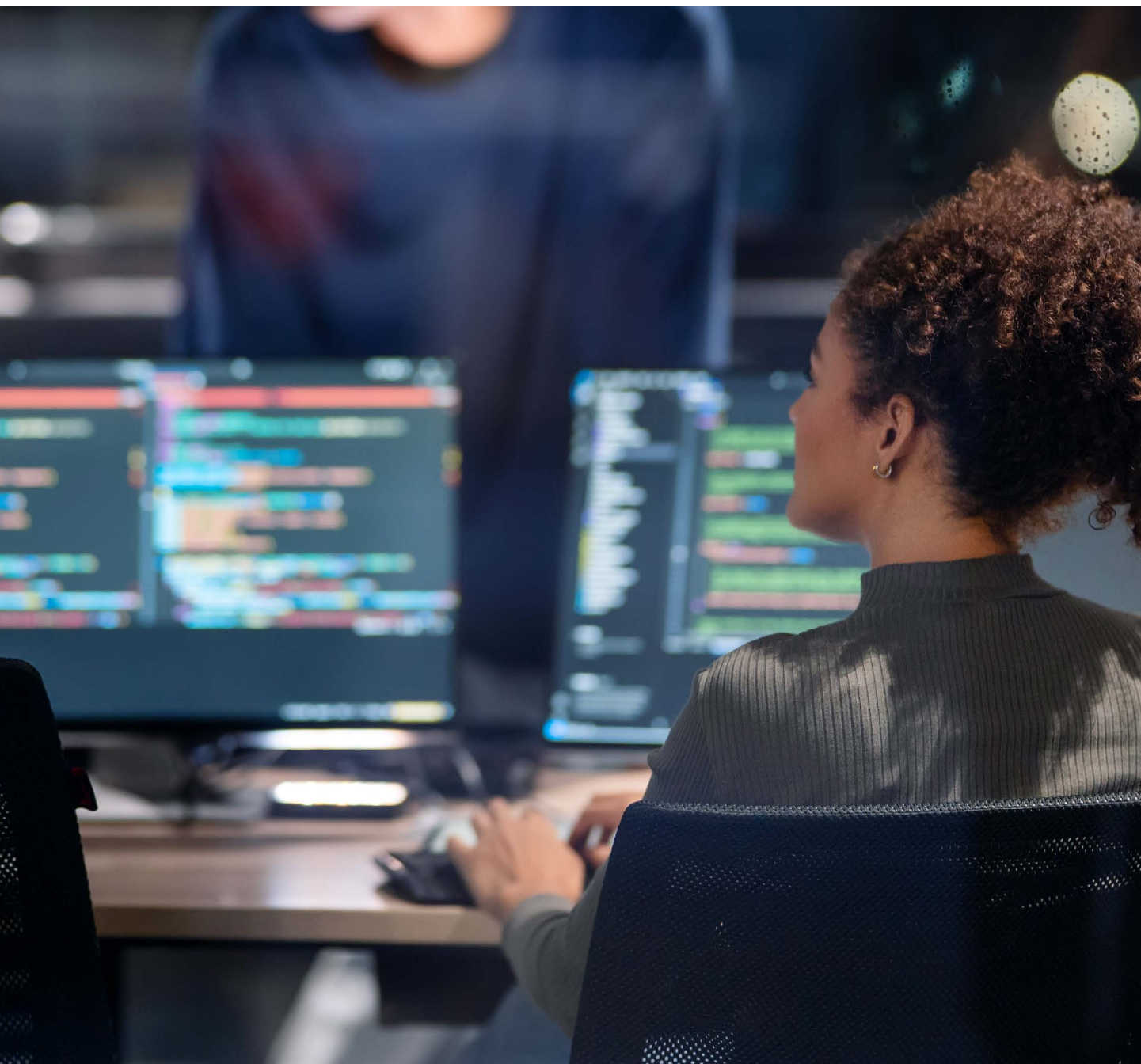
- What to do when data is missing
- What to do when signals conflict
- When to abstain and escalate

By assuming a non-ideal, degraded scenario, you can avoid the common problem of AI forcing a sub-optimal output. In many cases, it's better for AI not to engage at all rather than provide an incorrect output.





The AI-powered SOC in practice



Responding to UEBA alert triggers with AI, humans, and automation

User and entity behavior analytics (UEBA) establishes a baseline for normal user and device behavior, then monitors actions to provide alerts in case of suspicious activity. However, manually monitoring all the data points that signify malicious activity can lead to false positives and notification fatigue.

A mix of automation and human analysts – empowered by AI – can quickly identify the signs of bad actors or compromised systems without overloading teams.

During the **enrichment** phase, automation takes care of correlation checks, querying sign-in logs, and checking IP reputation and TOR or proxy use. Meanwhile, analysts are empowered with AI to identify sign-in trends, such as common IP addresses, login times and locations, known devices and browsers, previous logins, and whether multifactor authentication was used.

A disposition interpreter code block using LLMs and human effort can recommend a **disposition** based on investigation runbook logic.

Automation handles the majority of the **response** stage; it marks user risk, disables the account if necessary, and then revokes the session.

And finally, an LLM creates a **summary** in markdown, guided by a human analyst.

The result is a faster, more consistent UEBA process that places far less burden on limited analyst time and resources.

Phishing campaign detection

Whenever a user sends in a suspicious email, analysts must investigate whether it's a legitimate phishing attempt or a false positive. But with large amounts of unstructured data to look through, this can be a time-consuming task that often doesn't lead to meaningful action.

AI, automation, and human analysis also play a joint role in rapidly detecting and reviewing potential phishing campaigns.

During the **enrichment** phase, automation queries proxy and email logs to determine reach, checks IP reputation against threat intelligence, and submits the URL. Meanwhile, analysts are empowered with AI to analyze the head and body of the message, review existing anti-spam and anti-malware policies, and assess which anti-malware, anti-phishing, and safe link and attachment policies apply to specific users.

A disposition interpreter block of code can run using LLMs and human effort to recommend a **disposition** based on investigation runbook logic.

Humans and AI determine and execute an appropriate **response** based on the investigation findings – whether that involves blocking senders, pulling mail from users, or isolating hosts.

All these outputs are handed off to an LLM to create a **summary** in markdown, guided by a human analyst.

By combining AI and automation, you can save significant analyst time while accurately identifying legitimate phishing campaigns.



The AI-powered SOC: A Capgemini solution, leveraging Microsoft Security

Adopting the AI-Powered SOC raises the bar and increases resilience

Organizations achieve business resilience through cost-effective, outcome-based SOC services while eliminating threats, reducing attack surfaces, and continuously optimizing operations.

AI does not make a bad SOC better; it makes a disciplined SOC scale.

					
					
A balanced design yields results	Compliance	Productivity	Process	Resilience	
	• Monitor AI deployed for the SOC • Track SOC AI maturity • Use best practices from vendors providing tools to secure AI	• Intelligent triage • Workload reduction • Efficiency in the operating model • L1 activities evolve or disappear	• Quality detection • Simplify processes and redesign for AI • Less waste in investigation and response • Extend capabilities of analysts at all level	• Context rich insights • Alert quality goes up • Speed of decisions improves • Visibility extends into new sub-domains	





Capgemini, part of the Open Group, offers AI-powered SOC solution built on Microsoft technologies that helps you identify and deploy efficient, well-governed AI use cases in your SOC.

Capgemini's Modern Security Operations solution, integrated by Microsoft Sentinel, Microsoft Defender XDR, and Microsoft Entra ID, transforms SOC's into AI-enabled defense hubs.

It offers five key enablers to elevate every part of your SOC:

- **Standardized foundations** transform tooling and deliver data consistency to support AI use cases.
- **Detection and automation reviews** tune and optimize workflows while deploying automations that support AI detection, investigation, and response.

- **Cost-controlled investigations** drive low-cost, repeatable workflows for threat detection and hunting using promptbooks and AI-based skills.
- **Optimized implementations**, where AI deliverables are refined and use cases are built based on a prioritized backlog.
- **Continuous improvements** to keep pace with new threats and maintain readiness.

Capgemini's Modern Security Operations solution leverages the Microsoft Security ecosystem to transform cybersecurity operations, streamlining threat detection, triage, investigation, incident response, and continuous improvement across the SOC.

Begin your AI-powered SOC transformation today

The threats you face are rapidly advancing. Your SOC should be too.

The use cases we've explored in this guide represent a fraction of the threat scenarios an AI-powered SOC is designed to handle.

The reality is that each organization will have different environments, tools, risk profiles, and regulatory context. Capgemini can work with you to design a security operations model that matches your specific reality, and unlocks powerful new AI use cases based on your unique needs. We can also provide managed security services that scale with you.

To book a SOC assessment to establish your baseline and start your AI journey, get in touch.

cybersecurity.in@capgemini.com





Appendix: AI use case assessment prompt

SOC Use Case Suitability Evaluation

Microsoft Security Copilot – Governed & Cost-Bound Prompt

This prompt is intended for use with **Microsoft Security Copilot** in a **Security Operations Center (SOC)** to evaluate whether (insert task, e.g. vulnerability management) workflows should be handled by:

- Traditional automation
- Security Copilot
- A hybrid approach (automation first, AI with gates)

The goal is to prioritize **operational efficiency, cost control, repeatability, and auditability**.

Evaluation Criteria

Assess each use case using the four dimensions below.

1. Determinism

Is the task rules-based, repeatable, and reliably executable via scripts, playbooks, or workflows without interpretation?

- **High** → Favor automation
- **Low** → Favor AI or human judgment

2. Judgment

Does the task require synthesis, contextual reasoning, prioritization, or explanation across unstructured or semi-structured data?

- **High** → Favor Security Copilot

3. Auditability

Can inputs, prompts, reasoning, and outputs be made traceable, consistent, and repeatable for audit, justification, and SOC review?

- **Low auditability** → Do not recommend AI inline

4. Frequency & Cost

How often would this task run in normal SOC operations, and what is the impact of repeated AI invocation?

- **High frequency + deterministic** → Avoid AI
- **Low frequency or gated** → Acceptable for AI

Decision Logic (Mandatory)

- **High frequency + high determinism** → Automate (not AI)
- **High judgment + low frequency** → Security Copilot
- **High frequency + some judgment** → Hybrid with strong gates
- Never recommend AI if automation delivers the same result at lower cost.

Invocation Guardrails

Invocation Limits

- Do not recommend AI invocation:
 - More than **once per asset per vulnerability per review cycle**
 - More than **once per ticket** unless context materially changes
- Prefer **output reuse, caching, or memorization**.

Retry Rules

- If confidence is low or data is insufficient:
 - Allow **one retry only**, using the same prompt template
 - Do not chain exploratory follow-up prompts
- If uncertainty remains, **escalate to human review**



Gating Rules

- Automation must always execute **before** AI in hybrid workflows
- AI may only be invoked when:
 - Deterministic logic fails, or
 - Narrative explanation or justification is required
- AI outputs are **advisory**, not authoritative

Output Constraints

- Provide structured, concise reasoning
- Avoid speculative language
- Clearly separate **facts**, **assumptions**, and **interpretive judgment**

Token & Cost Budget

Token Budget

- **Soft limit**: ≤ 1,000 tokens per invocation
- **Hard constraint**: One inference per decision point
- Avoid multi-step reasoning chains unless explicitly required

Cost Controls

- Treat this workflow as **Medium-Cost AI**
- Single, well-scoped invocation
- Output must be reusable across tickets, dashboards, or reports
- Do not invoke AI:

- Per alert
- Per scan result
- Per CVE without aggregation or gating

Overrun Handling

- If token or cost limits would be exceeded:
 - Return a summary stating **AI is not cost-effective**
 - Recommend automation or human review

Task Instructions

For each use case:

1. Rate **Determinism**, **Judgment**, **Auditability**, and **Frequency** as:
 - Low / Medium / High
2. Recommend one execution model:
 - Automate
 - Security Copilot
 - Hybrid (Automation + AI with gates)
3. Prioritize which workflows should be enabled first, with a short justification based on SOC value and cost discipline

Use Cases [insert your use cases here]

1. Exploitability Context Enrichment
2. Vulnerability → Identity Exposure Mapping

3. Vulnerability → Internet Exposure Assessment
4. Patch Priority Scoring with Business Context
5. KEV Alignment & Risk Justification
6. Lateral Movement Risk from Vulnerability
7. Vulnerability → Detection Coverage Gap Analysis
8. False Positive Vulnerability Validation
9. Vulnerability Recurrence Root Cause Analysis
10. Executive Vulnerability Exposure Brief

Final Principle

Recommend where to use Security Copilot where it **improves decision quality**, not where it replaces reliable automation.

Favor **predictable, governable, and cost-aware** solutions over exploratory or conversational use.

Author



Mona Ghadiri

Vice President, Global Offer Lead for Cybersecurity Defense

Cloud Infrastructure Services

mona.ghadiri@capgemini.com

For more details, contact:

cybersecurity.in@capgemini.com

About Capgemini

Capgemini is the business transformation partner for enterprises in the age of AI. We help organizations imagine and build an intelligent, sustainable future, combining AI, technology and human ingenuity to transform how they operate, innovate and grow. With unique end-to-end capabilities spanning strategy, technology, engineering and intelligent operations, we bring together deep industry expertise and market-leading capabilities in AI, cloud and data to turn ambition into measurable business outcomes at scale. Supported by a robust ecosystem of partners and nearly 60 years of expertise, Capgemini is a responsible and diverse global organization of over 410,000 team members in more than 50 countries. The Group reported 2025 revenues of €22.5 billion.

Make it real.

www.capgemini.com

