



RETHINKING THE SOC:

From Linear Pipelines to Autonomous Systems

www.stellarcyber.ai | sales@stellarcyber.ai

Executive Summary

Security Operations Centers (SOCs) have reached an inflection point. Traditional, linear workflows—optimized for log ingestion, rule-based alerting, and manual triage—are collapsing under the weight of today's threat landscape. **The modern SOC must evolve into a recursive, intelligent system: the Autonomous SOC.**

This whitepaper introduces a purpose built approach to SOC transformation, blending verdict-first AI, structured validation known as Verdict Signal Checks (VSCs), a Stellar Cyber coined concept and framework, a modular validation framework that ensures alerts earn promotion through lightweight checks, human-in-the-loop design, and cost-effective automation. The result is a security operations architecture that doesn't just reduce noise; it redefines how decisions are made, where human input is needed, and how AI agents scale across the detection-response lifecycle.

Most importantly, this transformation is already happening. Early adopter customers testing autonomous case analysis and triage capabilities have observed a fundamental shift in the junior SOC analyst role. Instead of conducting detailed investigations, junior analysts are increasingly coordinating and passing along AI-generated case summaries. This is not a vision of the future – it is a change underway right now. As the technology matures, we expect similar evolution for Tier 2 and Tier 3 analysts, with their focus moving toward tuning, strategy, and proactive defense.

Table of Contents

- 1. From Linear to Recursive: Rethinking Security Workflows..... 3
- 2. Verdict-First AI: The Role of VSCs 5
- 3. Evaluation Sample Sets: Matching the VSCs 7
- 4. Human-Augmented SOC: The Pyramid of Influence 7
- 5. Economics of Security AI: Why LLMs Alone Don't Scale..... 9
- 6. Platform Architecture: Engineering for Security 11
- 7. The Shifting Role of the SOC Analyst.....12

1. From Linear to Recursive: Rethinking Security Workflows

For over 20 years, SOC teams have followed a predictable sequence: data ingestion → rule-based alert → manual triage → playbook response. While functional, this model suffers from three systemic weaknesses:

- **Noise Overload:** Rules fire without context, burying analysts in false positives.
- **Siloed Investigations:** Alerts are triaged in isolation, ignoring entity relationships.
- **Slow Response:** Manual steps stretch investigation and containment timelines.

The Autonomous SOC introduces a **recursive model**, where every investigative step can generate new questions, fetch new context, and inform next actions. What makes it recursive is the ability to pull in just-in-time data and context that were not part of the original alert payload—enrichments, intelligence lookups, or entity baselines brought in dynamically as the investigation unfolds. For example, an impossible travel alert might trigger the system to fetch ASN risk information, device compliance status, or cross-check login history across cloud services. This just-in-time enrichment illustrates how investigations evolve dynamically beyond the initial alert data.





During the linear workflow era, AI was introduced as a feature to streamline specific steps: ML-based detection for anomalies, Graph ML for cross-alert correlation, and AI investigator assistants to generate search queries. These additions reduced friction but remained incremental. AI was still just a feature of the linear workflow, not a paradigm shift.

The move to **Agentic AI** is more disruptive. Instead of assisting individual steps, AI agents become the **operating paradigm of the SOC**. They work recursively—reasoning, validating, correlating, and escalating—acting more like analysts than static tools. Instead of linear playbooks, investigations become **dynamic loops** that converge toward validated outcomes.

Human-Augmented Autonomous SOC

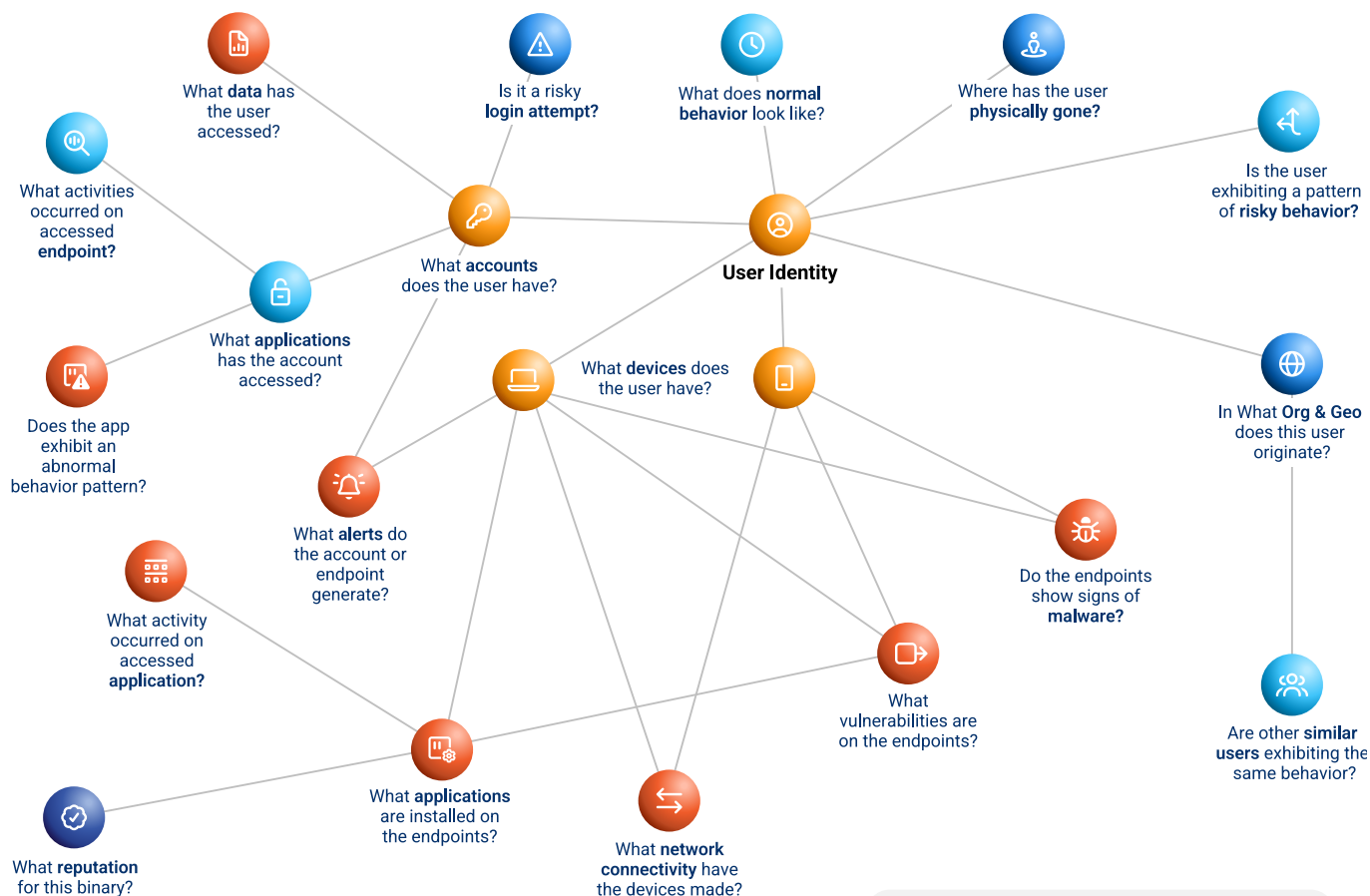


2. Verdict-First AI: The Role of VSCs

Traditional detection systems are **fire-first**: alerts are triggered by rules or statistical models and then pushed to humans to validate. This approach guarantees high noise and erodes analyst trust.

The **Verdict Signal Check (VSC)** framework reverses this model. Instead of asking analysts to disprove noise, VSCs demand that detections **earn their way up the confidence**

ladder through structured, lightweight validation. Each VSC is an explicit, modular check that simulates what an analyst would do in the first few minutes of triage. The visualization below illustrates how VSCs organize checks across observables, origins, impacts, and meta-analysis to drive verdicts.



VERDICT SIGNAL CHECK

- **Observable**
- **Origin** (What Happened?)
- **Impact** (So What?)
- **Meta** (History/Feedback/Prevalence)
- **Threat Level** (Just in Time TI Lookup)

Categories of VSCs

The design framework spans several domains:



Origin checks: Validate whether the observable itself looks suspicious. Examples: unknown process names, rare parent-child lineage, login from unfamiliar IP ranges.



Impact checks: Assess what the event could lead to if malicious. Examples: sensitive service access, unusual privilege escalations, data staging activity.



Meta-analysis: Compare against historical prevalence and prior verdicts. Example: "Has this host or user triggered similar alerts before, and were they previously dismissed or confirmed?"



External context: Just-in-time enrichment from threat intelligence (VirusTotal, AbuseIPDB, IPInfo, MaxMind Pro, etc.).

Examples in practice:



Endpoint triage: VSCs examine suspicious process behaviors. Legitimate system binaries or processes in trusted directories are down-scored, while temporary executables with encrypt arguments are escalated.



Identity triage: VSCs validate impossible travel, risky ASN logins, and device compliance. A login from a managed endpoint previously associated with the user might be suppressed, while a first-time login from a non-compliant device is elevated.



Threat intelligence triage: VSCs weigh the credibility of IP/domain risk. IPs flagged across multiple feeds (VirusTotal, AbuseIPDB) are weighted higher than a single-source flag.

Verdict Signal Check Examples

VSC Name	Primary Question (VSC #1)	Follow-up Question (VSC #2)	Possible Outcomes/Reasoning
Process Prevalence Check	Has this process hash been observed running across more than 20 different hosts in the last 24 hours?	Is the process hash signed by a trusted publisher or present in a known-good catalog?	If yes → benign software (FP). If no → malware propagation (TP).
Password Change Timing	Has the user's password been changed within 60 minutes of this login failure?	Did both events originate from the same device/IP?	Same device/IP → benign FP. Different IP → suspicious TP candidate.
Past Analyst Verdicts	Were similar alerts on this host overridden as false positives by analysts in the past week?	Has the context changed (binary version, parent process, host role)?	No change → suppress (FP). Changed context → escalate (TP).

Dynamic Signal Strength

Each VSC carries a **signal strength** (low, medium, high) and an **execution mode** (auto or manual assist). For instance:

- **Meta-Previous_Override:** A prior analyst verdict is treated as high-strength auto evidence.
- **Meta-Host_Prevalence:** A host repeatedly raising the same alert type is a high-strength supporting signal.
- **Meta-Previous_Verdict:** Similar alerts previously resolved as false positives may lower the weight (medium-strength signal).

This design ensures verdicts are not the output of one detection but the **accumulated reasoning of multiple checks**, each traceable back to its source.

Why VSCs Matter



Explainability: Every promoted alert comes with a rationale: “Escalated due to impossible travel validated against geolocation and ASN risk.”

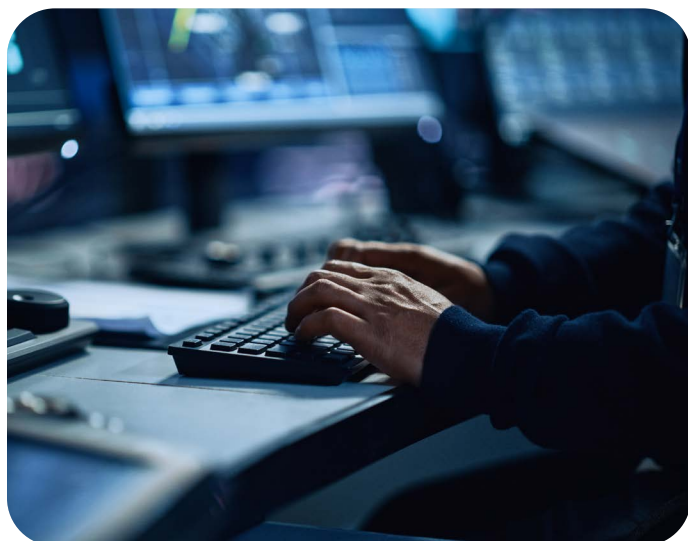


Efficiency: Many noise alerts die in the VSC layer, never burdening analysts.



Trust: Analysts gain confidence that automation is not opaque, but reasoned and auditable.

VSCs form the **bridge between raw detections and trusted verdicts**. They are the “mini-investigations” that allow AI to act like a Tier 1 or Tier 2 analyst, but at machine scale.



3. Evaluation Sample Sets: Matching the VSCs

For VSCs to succeed, AI agents must be trained and evaluated on **sample sets that reflect operational reality**. Precision and recall scores alone are inadequate if they ignore adversarial edge cases or benign-but-noisy conditions.

A comprehensive evaluation framework must:

- Include **clean, noisy, and adversarial** samples
- Cover diverse environments (cloud, on-prem, mobile, IoT)
- Represent **rare, high-impact events** alongside common benign behaviors
- Account for **data freshness issues** (delayed logs, missing enrichment)

The goal is to measure not just model accuracy, but **decision usefulness in a SOC context**. Evaluation must bridge detection engineering, AI development, and analyst feedback—ensuring that every promoted verdict is both technically correct and operationally valuable.

4. Human-Augmented SOC: The Pyramid of Influence

The skill set required for SOC analysts is shifting toward becoming **AI supervisors**. In the human-augmented Autonomous SOC, analysts validate, approve, confirm, or override AI verdicts and decisions. Their influence –and ultimately their performance– depends on the type and depth of their feedback. A more senior analyst will have and should have a greater impact on the AI. The hierarchy of responsibility can be measured by the level of influence humans exert over AI agents. This observation is likely true in many AI-driven domains, but it is particularly evident in SOC operations. It is a fascinating topic that goes beyond the scope of this whitepaper.

Automation is powerful; analyst oversight is indispensable. To make them complementary, we build upon the concept of a **Pyramid of Influence**, refined with feedback impact tiers as described in Stellar Cyber’s blog post “*From Pyramid of Pain to Pyramid of Influence*.” This refined pyramid clarifies the types of analyst feedback, how that feedback influences the system, and where human input is non-negotiable.

Feedback Impact Tiers

At the base are simple tags; at the apex are strategic changes. Each tier represents feedback that has increasing influence:

1. **Passive Classifications** (Lowest influence): "True Positive/False Positive" flags without additional context. Useful for measuring precision but limited in driving change.
2. **Contextual Annotations**: FP/TP feedback enriched with descriptive notes: what exactly went wrong (process, host, schedule). These feed directly into detection tuning and alert suppression rules.
3. **Overrides & Suppressions**: Strong signals from analysts that certain detection logic be disabled or adjusted (e.g., known benign behavior, trusted devices). These modify thresholds, VSC weights, and suppress noise at the source.
4. **Detection Innovation & Logic Changes** (Highest influence): Analysts propose new detection features, VSC categories, or model adjustments. System logic evolves in response.

When Human vs Agent Acts

- **Autonomous Tier**: The system acts without oversight when detections have passed multiple VSCs with high confidence and there is no contradicting analyst feedback history.
- **Collaborative Tier**: The system generates alerts or verdicts but includes analysts in the decision loop—especially when signal strength is borderline, or when novel behavior is detected.
- **Judgment Tier**: Human analysts step in for rare, sensitive, or ambiguous cases: insider threat, strategic decisions, policy exceptions.

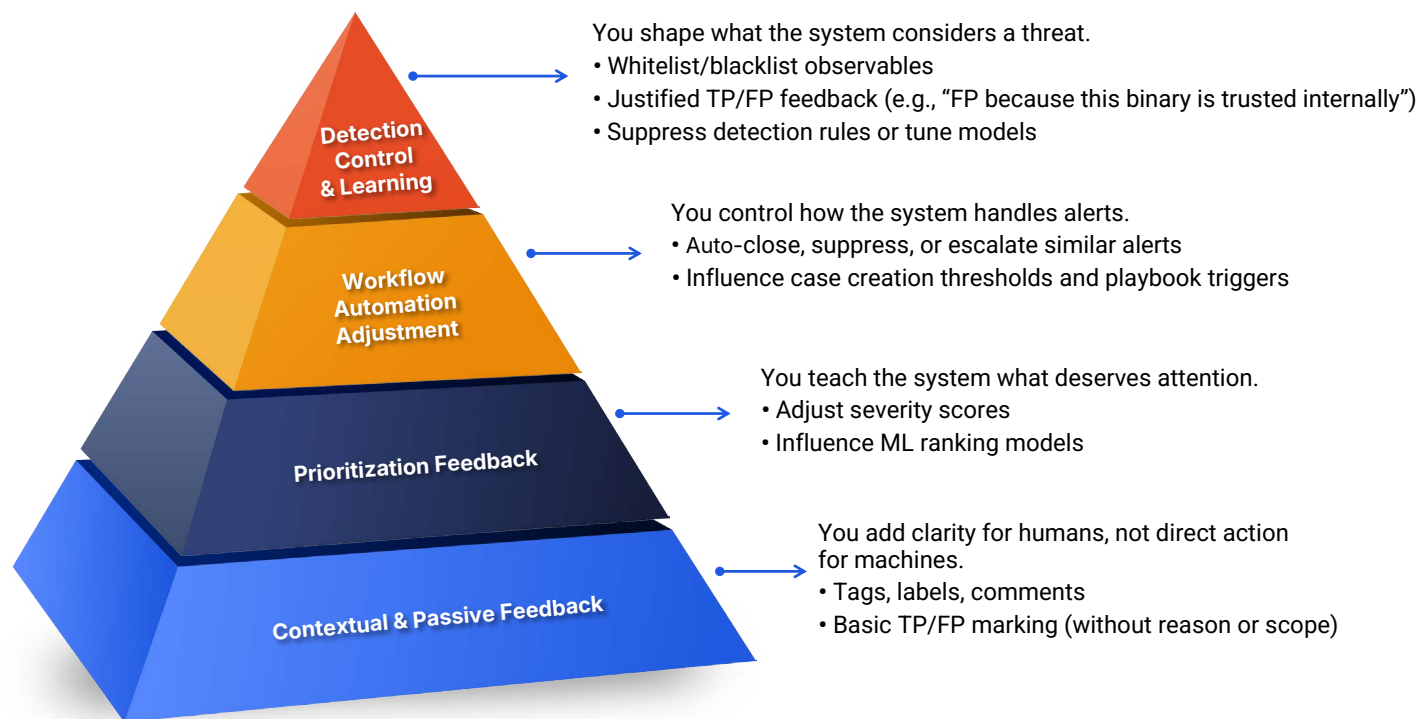
Feedback Loop: How Analyst Input Transforms the System

1. **Capture**: Analyst submits feedback via the UI: tagging, annotating, or overriding system output.
2. **Classify & Weight**: Feedback is categorized into the impact tier. The system logs strength, context, and patterns.
3. **Action**: Depending on tier, action may include: suppressing future alerts, adjusting model weights, updating VSC configurations, or deploying new detection logic.
4. **Measure**: System tracks whether changes reduce noise, improve precision, or accelerate response time. This becomes input back into evaluation.
5. **Trust Reinforcement**: Results visible to analysts (e.g., "Your suppression prevented 1,200 false alerts") build trust and incentive.

"At Stellar Cyber, we believe visibility drives action. By unifying data across every source, our platform transforms how security teams work — turning hours of analysis into minutes of insight. When organizations can see everything clearly, they act with confidence, speed, and intelligence. That's the true power of open, connected cybersecurity."

— Aimei Wei, CTO and Founder

The following illustration from Stellar Cyber's blog post *"From Pyramid of Pain to Pyramid of Influence"* can be used to visualize this model:



5. Economics of Security AI: Why LLMs Alone Don't Scale

In the current market, LLMs are often marketed as a solve-it-all solution for the SOC. Many vendors and startups position them as engines that can replace existing detection, correlation, and triage pipelines. The reality is more nuanced. While the hype suggests LLMs can take over end-to-end operations, the economic and technical realities show that they are best used as assistants, not engines.

LLMs excel at **creative acceleration** –writing regex, building parsers, drafting correlation logic, or summarizing complex cases– but they are not designed for real-time, scalable detection. Think of them as the code whisperer, not the pipeline processor.

Where LLMs Shine:

- Generating regex, parsers, or correlation logic.
- Summarizing alerts, cases, or incident timelines.
- Explaining threat scenarios in plain language for analysts or executives.

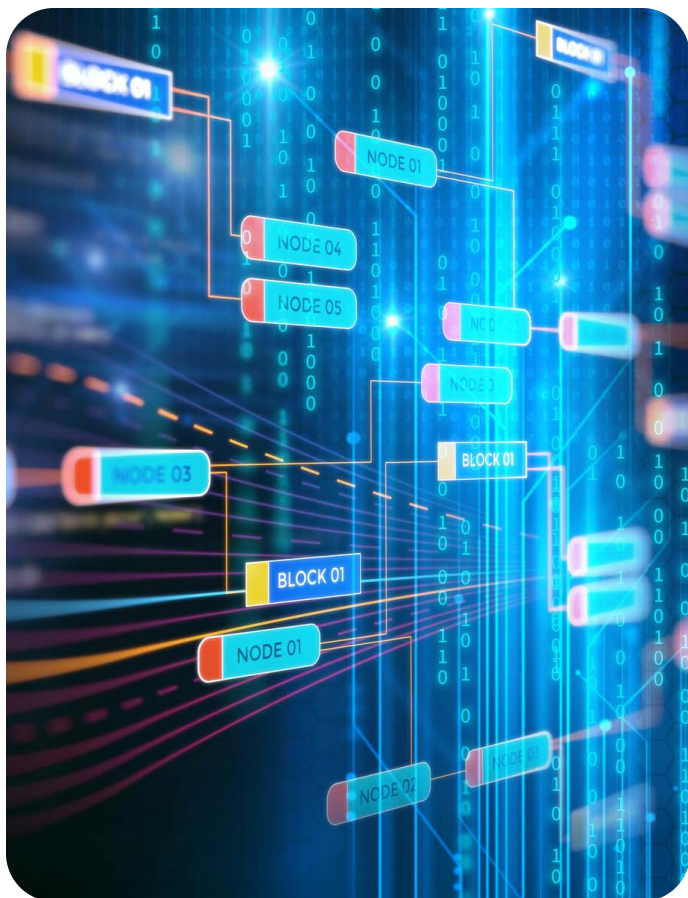
Where Traditional ML Excels

- Real-time detection and streaming analytics.
- Scalable anomaly detection (UEBA, DNS tunneling, credential stuffing).
- Low-latency, low-cost classification at event scale.

Economic Realities and the Cost Equation

Even with dramatic growth, LLMs remain cost-prohibitive for mass ingestion or real-time alert triage:

- **Token Limits:** Context windows have expanded 32× (from ~4K in 2023 to ~128K today).
- **Token Speed:** Throughput has improved ~1,500×, now ~94K tokens/sec (~250 events/sec).
- **Cost:** Despite efficiency gains, running LLMs at SOC scale would still cost tens or hundreds of thousands of dollars per month.



Additionally, in AI-agent use cases such as **Verdict Signal Checks (VSCs)**, there may be numerous back-and-forth calls and collaborations between agents. This recursive reasoning is powerful but comes with a financial risk: unresolved loops could lead to uncontrolled token consumption. To make this sustainable, **cost boundaries and safeguards** must be built in to prevent runaway expenses.

In other words: LLMs now handle deep context and summarization with remarkable efficiency, but they remain unsuitable for streaming detection, bulk log parsing, or alert-scale classification. They augment the SOC—they do not replace it.

Practical Guidance

- **Let GenAI generate the parser.**
Don't run 20 million logs through it.
- **Let it explain a threat scenario.**
Don't expect it to monitor real-time streams.
- **Use it for triage & summarization.**
Not for mass-scale ingestion.

The right model at the right stage is key:

- **Rules:** Fast, interpretable, cost-efficient.
- **Traditional ML:** Domain-specific, scalable, reliable.
- **LLMs:** Contextual reasoning, summarization, creative assistance.

This hybrid approach ensures that AI is **economically sustainable, technically viable, and operationally effective**.

Cost-Control Framework

To bound costs and avoid runaway token consumption in recursive agent workflows, a cost-control framework should be implemented:

- **Recursion Depth Limits:** Restrict how many times agents can call each other before escalation or termination.
- **Budget Caps:** Define token or cost thresholds per case, alert, or investigation.
- **Fallback Paths:** When budget is exceeded, revert to deterministic rules or traditional ML for resolution.
- **Monitoring & Auditing:** Track token usage, agent interactions, and financial impact as part of SOC observability.

These safeguards ensure that the benefits of agentic reasoning are realized without incurring uncontrolled costs.

Illustrations

To visualize these trade-offs and improvements, include the following figures:

- **LLMs Are Powerful – But Not Built for All Tasks**
- **LLM Power: Growth, Not Replacement (Yet)**
- **Cost Considerations of LLMs at Scale**

6. Platform Architecture: Engineering for Security

A common question arises: how is this different from simply copying and pasting alerts into a generic tool like Microsoft Copilot to get an explanation? The answer lies in the architecture. Our pipeline is purpose-built for security operations, not productivity assistance.

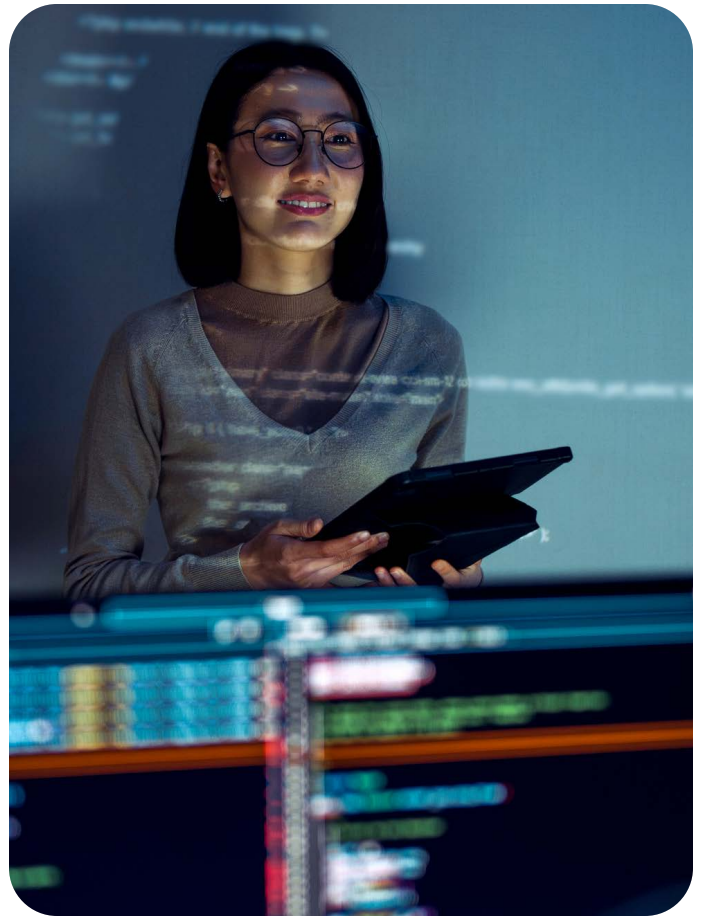
- **Field-Aware Processing:** Raw JSON is tailored to focus on the entities that matter most—users, IPs, hosts, and timestamps.
- **Scalable Case Handling:** The system processes dozens or even hundreds of alerts within a case without losing coherence or storyline.
- **Precise Sequencing:** Raw UNIX timestamps are processed to ensure the timeline reflects exactly how the attack unfolded.
- **Noise Reduction via Graph Models:** The system addresses the needle-in-a-haystack problem by using graph analysis to reduce noise before summarization.

This is the difference between a productivity copilot and a **purpose-built SOC pipeline**: one provides ad hoc explanations, while the other delivers **trusted, repeatable case analysis at scale**.

To achieve this, the Autonomous SOC must also follow a set of engineering principles designed for security rather than general business AI platforms:

- **Queue-Aware AI:** Agents prioritize based on backlog, case severity, and analyst availability.
- **Cost-Aware Processing:** Lightweight checks are applied early, with LLM power reserved for high-value cases.
- **Entity-Centric Design:** All context (users, hosts, files, IPs) is unified into case views.
- **Explainability by Default:** Every decision, score, and action comes with an auditable rationale.

Together, these characteristics ensure that the Autonomous SOC pipeline is not only autonomous but also **transparent, accountable, and operationally aligned** with the realities of security operations.



"At Stellar Cyber, our mission is to deliver intelligence that drives action. By integrating all critical data sources into a single open platform, we empower SecOps teams to uncover insights in minutes instead of hours. Simplicity, speed, and visibility are built into everything we design — transforming complexity into clarity for every customer."

– Subo Guha, SVP, Product Management

7. The Shifting Role of the SOC Analyst

The introduction of verdict-first AI, VSCs, and human-agent collaboration fundamentally changes the role of the SOC analyst. One customer testing this capability observed that the system was able to perform deep case analysis during a simulated web application penetration test. What previously would have required hours of manual log review and evidence stitching was automatically summarized into a comprehensive case narrative with clear verdicts. The analyst’s role shifted from line-by-line investigation to **reviewing, approving, and contextualizing** the system’s findings.

This represents a broader industry trend. The matrix below contrasts how analyst roles are shifting, and highlights the new skills required at each tier:

Analyst Tier	Traditional SOC Role	Autonomous SOC Role	New Skills Required
Tier 1 (Junior)	Manual alert triage, filtering noise, escalation	AI supervisor, enrichment, process coordination	Understanding AI outputs, validating/overriding AI decisions, basic scripting for automation, communication & coordination skills
Tier 2 (Intermediate)	Case investigations, alert correlation	Guiding AI feedback loops, refining detection logic, contextualizing cases	Detection engineering, query design, feedback annotation, cross-domain knowledge (cloud, identity)
Tier 3 (Senior)	Escalations, advanced investigations	Strategic analysis, threat hunting, purple teaming, evolving detection/VSC design	Threat hunting, adversary simulation, model/VSC design, mentoring, strategic oversight

In this future, the SOC is not diminished—it is **elevated**. Analysts at all levels move from reactive alert triage to higher-value contributions: tuning, oversight, strategy, and proactive defense.



Conclusion: The Path Forward

The Autonomous SOC is not about replacing analysts—it is about **augmenting them with trusted, recursive, and explainable intelligence**. By combining VSC-driven validation, comprehensive evaluation datasets, a refined human-agent Pyramid of Influence, and a hybrid AI strategy, we can:

- Silence noise and restore analyst focus
- Scale decision-making without exponential cost
- Deliver auditable, human-readable rationales for every action
- Create a SOC that adapts, learns, and improves with time

What we are observing with our early adopter customers confirms this transformation: the **junior SOC analyst role is actively evolving today**. Instead of performing detailed investigations, they are increasingly coordinating

and passing along the AI's comprehensive summaries. This shift is not a prediction—it is happening now. Early adopters are leveraging autonomous case analysis and auto-triage in production, and the junior analyst role is already transitioning into a more supervisory and high-value coordination function.

To thrive in this new paradigm, SOC analysts at all levels must develop **new supervisory skills**: understanding AI outputs, knowing the type of influence they can exert on the model, and learning how to steer AI decisions effectively. These skills are essential for maximizing both human and AI performance in the SOC.

The future of security operations is **recursive, validated, and human-centered**. It is not linear pipelines or dashboard fatigue—it is dynamic loops of trust and collaboration.

By shining a bright light on the darkest corners of security operations, Stellar Cyber empowers organizations to see incoming attacks, know how to fight them, and act decisively – protecting what matters most. Stellar Cyber's award-winning open security operations platform includes NG SIEM, NDR, Open XDR, and Multi-Layer AI™ under one license. With almost ⅓ of the top 250 MSSPs and over 14,000 customers worldwide, Stellar Cyber is one of the most trusted leaders in security operations. Learn more at <https://stellarcyber.ai>.