

CORE FUNCTION BRIEF

Generative and Agentic AI

A core function of the ClearSkies iISOC platform

Two AI modalities, one governance gate, drafting and investigating on behalf of the function that owns the decision.

Contents

1	Executive Summary
----------	--------------------------

2	Why Manual Drafting and Investigation Fall Short
	Narrative work consumes analyst time
	Investigation depth capped by available hours
	Draft-to-decision lag
	Improvement that depends on memory
	The cost of the workaround

3	How Generative and Agentic AI Works
	The five stages
	How an output is judged

4	Generative and Agentic AI in the Platform
	What it draws on, what it writes back, and what it does not do
	A worked example
	Delivery across many tenants

5	ATT&CK Alignment and AI Governance
----------	---

6	Questions Raised in Evaluation
----------	---------------------------------------

7	The Benefits
	At a glance
	What changes, and for whom
	What sets it apart

8	Summary
----------	----------------

SECTION 1

Executive Summary

IN ONE SENTENCE.

Generative and Agentic AI is the core function that drafts, investigates and validates on behalf of every other core function that needs it, generative for narrative and explanation, agentic for correlation and candidate action, with every output passing the requesting function's own governance gate before it is used.

Drafting and investigating consume analyst time in a conventional security operation. Writing an incident summary takes nearly as long as the investigation itself. Correlating weak signals by hand is capped by the hours available rather than by what is worth investigating. And a candidate detection or response step that exists in someone's head still has to be typed, tested and reviewed before it helps anyone. None of this scales with alert volume, and none of it improves unless someone remembers to change the process next time.

Generative and Agentic AI answers this with two modalities behind one governed loop. Generative AI turns normalized context into a narrative, a summary, or a draft artifact. Agentic AI investigates, correlates and validates evidence, producing a candidate finding or a candidate action for review. Neither modality decides anything on its own authority: every output passes the review gate the requesting function defines before it reaches production, and the outcome of that review feeds back to improve the next draft.

The same two modalities serve every core function that needs them: Detection Factory for authoring and validating detection content, SOAR for investigation and candidate response options, Threat Hunting for hypothesis generation, and KPIs and SLAs for report drafting. One capability, bounded task by bounded task, rather than a separate AI feature built inside each.

THE FUNCTION IS PRECISELY SCOPED.

Generative and Agentic AI drafts and investigates on request. It does not compute a risk score, which is the Risk Scoring Formula. It does not decide who works an alert, which is Intelligent Alert Routing. And it does not authorize or execute a response, which is SOAR and Playbooks. Each of those is a separate core function with its own brief, and none of them is described here as part of this one.

Two modalities

One governance gate

Four consuming functions

Generative for narrative, agentic for investigation

Nothing production-facing bypasses review

One capability, bounded task by bounded task

SECTION 2

Why Manual Drafting and Investigation Fall Short

Drafting and investigation are underinvested relative to the rest of the security operation, and the consequences compound quietly rather than surfacing in an incident review. Four failure modes explain why, and each is answered by a named stage of the loop in section 3.

2.1 Narrative work consumes analyst time

Writing an incident summary, a forensic timeline, or an executive report by hand takes time comparable to the investigation itself, and the work starts over from a blank page for every incident regardless of how similar it is to the last one.

2.2 Investigation depth capped by available hours

An analyst can correlate only so many weak signals manually before the return stops justifying the hours. Deep investigation is reserved for alerts that already look serious, so a subtle pattern that never quite looks serious enough goes uninvestigated.

2.3 Draft-to-decision lag

A candidate detection, a candidate response step, or a report exists as an idea before it exists as something reviewable. Turning the idea into a draft takes time a live incident does not have, so the gap between knowing what should happen and having something to approve is where urgency is lost.

2.4 Improvement that depends on memory

The lesson from a closed incident lives in a person's head or a retrospective document, not in a structure that changes how the next similar draft is written or the next similar signal is investigated. The same gap gets rediscovered rather than closed.

2.5 The cost of the workaround

The usual response is more senior analyst time allocated to writing and correlating rather than to deciding, which is the scarcest and most expensive hour in the operation. Cost scales with incident volume and with reporting obligations, while the underlying problem, drafting and investigating that do not compound into anything reusable, is untouched.

Industry research consistently ties a substantial share of analyst time to report writing and manual correlation rather than to decision-making, though the specific figure is not yet sourced for this brief.

SECTION 3

How Generative and Agentic AI Works

Generative and Agentic AI is one continuously operating loop native to the TDIR core, invoked task by task rather than running unbounded. A sibling core function submits a bounded request at one end, and a governed, auditable output returns to it at the other.

Figure 1. The five-stage governance loop.

3.1 The five stages

Each stage is named once here and referenced by name throughout the rest of this brief, paired with the failure mode from section 2 that it answers.

STAGE	WHAT HAPPENS	THE FAILURE MODE IT ANSWERS
Request	A sibling core function, Detection Factory, SOAR, Threat Hunting or KPIs and SLAs, submits a bounded task and the normalized context it requires.	Draft-to-decision lag
Generate	The generative modality drafts a narrative, translation, or candidate artifact, a summary, an explanation, or draft content, from that context.	Narrative work consumes analyst time
Investigate	The agentic modality investigates, correlates and validates evidence relevant to the task, producing a candidate finding or a candidate action for review.	Investigation depth capped by available hours
Govern	The generated or investigated output passes the review gate the requesting function defines, human or policy based, before it is used. Nothing production-facing skips this stage.	Draft-to-decision lag
Learn	The gate outcome, accepted, edited or rejected, feeds back to improve future drafts and investigations for that same task type.	Improvement that depends on memory

The two modalities are distinct in kind. Generative AI produces language: a narrative, a translation, a summary. Agentic AI produces a validated finding or a candidate action. Both are drafts until the Govern stage, and neither reaches production on its own authority.

3.2 How an output is judged

Quality is measured rather than assumed, on definitions a provider can report and a customer can audit. The measures below are reported natively and per task type. They are defined here because the same metric name is measured differently by different vendors, and target values are agreed per engagement rather than asserted.

MEASURE	THE DEFINITION USED	WHAT A POOR VALUE TRIGGERS
Gate acceptance rate	Outputs accepted without edits at the requesting function's gate, as a share of all outputs for that task type.	A review of the prompt, the context or the training signal for that task type.
Time to draft	Elapsed time from a task request to a reviewable output being ready for the gate.	Investigation of the specific stage introducing the latency.
Edit rate at gate	Outputs modified by the reviewer before use, as a share of all outputs for that task type.	A review of whether the edits reflect a pattern the Learn stage should absorb.
Investigation depth	Tasked alerts where agentic investigation added at least one corroborating or disconfirming signal beyond the original request.	A review of the context supplied with the task, which may be too narrow to work from.
Learning uptake	Gate decisions that measurably change draft quality for the same task type on subsequent runs.	A review of whether gate outcomes are being recorded in a form the Learn stage can use.

A rejected draft is a working gate, not a failed function. A task type with a nonzero rejection rate is not necessarily underperforming. The gate rejecting a candidate that does not meet the requesting function's bar, and the Learn stage absorbing why, is the mechanism protecting production quality working as intended.

SECTION 4

Generative and Agentic AI in the Platform

4.1 What it draws on, what it writes back, and what it does not do

What it draws on. Normalized telemetry and the shared entity model from the Centric-AI Fabric; the bounded task and its context from the requesting core function; and, for the Learn stage, the prior gate outcomes recorded for that task type. Each is an input to a specific stage rather than a general capability the function generates itself.

What it writes back. The gated draft, finding or candidate, together with its audit trail, published into the shared model and consumed by the requesting core function. The relationship is one way: nothing returned to this function changes an output already gated, and the Learn stage draws on the requesting function's own recorded outcome for that task rather than on a separate, untraceable score.

What it does not do. It does not compute a risk score, decide who works an alert, or authorize or execute a response. It holds no telemetry of its own, and it owns no governance gate other than the general checkpoint every output passes: the criteria a specific gate applies belong to the function that requested the task.

4.2 A worked example

The chain below is illustrative and technically representative rather than drawn from a named engagement. Detection Factory submits a Request: draft and validate candidate detection logic from a confirmed hunting finding. Generate drafts three candidates generalized to the underlying technique, each tagged with the ATT&CK identifier Detection Factory supplied. Investigate tests the candidates against ninety days of normalized telemetry, producing a precision estimate for each. Govern applies Detection Factory's own promotion criteria: two candidates clear the threshold and one does not. Detection Factory promotes the two into production. Learn records the rejected candidate's shortfall against the threshold, which sharpens the next draft Generate produces for a similar finding.

4.3 Delivery across many tenants

The same two modalities and the same five-stage loop serve every tenant. What changes per tenant is the context supplied with each task and the gate criteria the requesting function applies. What travels between tenants is the mechanism and the task-type learning it accumulates; what never travels is one tenant's telemetry, drafts or gate outcomes. A provider can therefore explain any drafted or investigated output to any customer, from the same mechanism, without a bespoke AI build per tenant.

Because the function is switched on with the platform rather than deployed separately, time to value is a matter of the first task types being exercised rather than an onboarding project. In the first week, requesting functions begin submitting tasks and every output passes a gate from the outset. Through

the first month, gate acceptance rate and edit rate accumulate enough history per task type to be read as figures rather than defaults. At steady state, learning uptake is tracked as an ongoing measure of whether the loop is closing.

SECTION 5

ATT&CK Alignment and AI Governance

Generative and Agentic AI does not map detection content and does not compute technique relevance itself. Its generative modality translates a technique identifier Detection Factory has already attached into a human-readable narrative for an analyst, and its agentic modality assists Detection Factory in drafting and testing a candidate mapping, which Detection Factory alone promotes.

GOVERNANCE IS THE SUBJECT, NOT A SIDE NOTE.

Because this function is the platform's AI capability itself, its governance obligations are stated directly rather than folded into a framework table. Every generated or investigated output passes a human or policy review gate, the audit trail records what was generated, what was reviewed and what was decided, and no sentence in this brief implies an output reaches production without that review.

FRAMEWORK OR STANDARD	ITS ROLE IN GENERATIVE AND AGENTIC AI
MITRE ATT&CK (Enterprise, Cloud and Identity)	Technique identifiers translated into narrative form, and drafting assistance for Detection Factory's own mapping. Neither computed nor promoted by this function.
EU AI Act	The governance obligations attaching to the generative and agentic layer specifically: the review gate, the audit trail and the human decision point named in section 4.1.
NIST Cybersecurity Framework 2.0	Outcome language used where a generated report discusses aggregate posture to an executive or board audience.

Evidence, not compliance. The platform states which obligations the governed loop produces evidence for. It never states that the presence of a review gate makes an organization compliant, and control-level mapping is held by the Regulatory Frameworks core function.

SECTION 6

Questions Raised in Evaluation

The questions below recur in buyer evaluations, each answered directly against the mechanism described in section 3.

QUESTION, AS A BUYER ASKS IT	ANSWER
Does Generative and Agentic AI decide anything on its own authority?	No. It drafts and investigates against a bounded task. The requesting function's gate decides whether the output is used.
Does it compute the risk score or decide who handles an alert?	No. Those are the Risk Scoring Formula and Intelligent Alert Routing. This function is not described as performing either.
Does it ever execute a response action directly?	No. It may investigate and draft a candidate action. SOAR authorizes a response and Playbooks executes it.
What stops a poor-quality draft from reaching production?	The governance gate the requesting function defines. No sentence in this brief, and nothing in the platform's design, allows an output to bypass it.
How does it improve over time?	The Learn stage draws on the requesting function's own recorded gate outcomes, accepted, edited or rejected, for that task type.
Is this the same thing as the AI-SecOps Autonomous Analysts?	No. This is a core function that drafts and investigates on request. The AI-SecOps Autonomous Analysts are a native add-on with their own brief, licensed separately and operating under their own authorized scopes.
Does the same capability behave differently per tenant?	The task context and the gate criteria are configured per tenant. The five-stage mechanism itself is the same across every tenant.

SECTION 7

The Benefits

The value of an AI capability in security operations is settled by what it removes from an analyst's day without removing anyone's judgment from the decision. The outcomes below are countable: how much of the drafting starts from context rather than a blank page, how much investigation happens that would otherwise not have been afforded, and whether every one of those outputs can be shown to have been reviewed.

7.1 At a glance

BENEFIT	WHAT PRODUCES IT
Drafting starts from context rather than a blank page	Generate turns the normalized context supplied with the task into a first draft for review.
Investigation depth stops being capped by available hours	Investigate correlates and validates evidence on every tasked alert, not only the ones that already look serious.
Nothing production-facing is used unreviewed	Govern applies the requesting function's own gate to every generated or investigated output.
Quality improves without anyone remembering to improve it	Learn absorbs the accepted, edited or rejected outcome for each task type and shapes the next draft.
AI governance is evidenced rather than asserted	The audit trail records what was generated, what was reviewed and what was decided, output by output.
One capability rather than an AI feature per function	The same two modalities serve Detection Factory, SOAR, Threat Hunting and KPIs and SLAs from one mechanism.

7.2 What changes, and for whom

A governed drafting and investigation capability changes what each role spends time on, and what each role can show for it afterwards.

ROLE	WHAT CHANGES
SOC analyst	Time goes to reviewing and validating a draft rather than producing it from a blank page, and the investigation arrives with corroborating evidence already gathered.
SOC manager	Drafting and investigation throughput becomes a reportable, auditable number rather than an assumption about how busy analysts are.
CISO and risk	Executive-ready narrative is available without a manual writing cycle, and AI governance is evidenced by an audit trail rather than asserted in a policy document.
Service provider practice lead	The same drafting and investigation capacity applies across every tenant, scaling with the customer base rather than with headcount dedicated to writing and correlating.

7.3 What sets it apart

- **Bounded by the requesting function.** Every task is scoped by the core function that owns the decision, so drafting and investigating never substitute for that function's own governance.
- **Nothing skips review.** One governance checkpoint applies to every generated or investigated output before it reaches production, whichever modality produced it.
- **Learning uses the outcome the requester already measures.** Closed-loop improvement draws on the requesting function's own quality decision rather than on a separate, untraceable score nobody can audit.

THE VALUE IN ONE LINE.

An ungoverned AI produces output faster than anyone can check it. A governed loop produces output that has already been checked, by the function that owns the decision, with the record to prove it.

SECTION 8

Summary

Generative and Agentic AI is how the platform drafts and investigates on behalf of the functions that decide. It is one core function among twelve, scoped to a five-stage governance loop that runs continuously inside the TDIR core rather than as an unbounded intelligence claiming every decision in the security operation.

- ▶ Generative AI drafts narrative and explanation; agentic AI investigates and validates. Neither decides on its own authority.
- ▶ Every task is bounded by the core function that requested it, and every output passes that function's own governance gate.
- ▶ Learning draws on the requesting function's own recorded outcome, not a separate, untraceable score.
- ▶ The same two modalities serve Detection Factory, SOAR, Threat Hunting and KPIs and SLAs from one mechanism.
- ▶ The function drafts and investigates, and leaves scoring, assignment and response authorization to the core functions built for those decisions.

Forward-looking statement: any statement about future capability is forward-looking and non-binding, and capability described in the present tense is capability available at the publication date, subject to the deployment configuration in force. This document describes drafting and investigation capability and does not constitute legal advice on any customer's obligations under the EU AI Act or any other instrument.