

CORE FUNCTION BRIEF

Threat Intelligence

A core function of the ClearSkies iISOC platform

Indicators and named threats, tagged, rated and republished as the evidence changes.

Contents

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

2	Why Unmanaged Feeds Fall Short
	Indicators without context
	Severity that is not consistently applied
	Global patterns invisible at the tenant level
	Advisories that go stale
	The cost of the workaround

3	How Threat Intelligence Works
	The five stages
	How a rating is judged

4	Threat Intelligence 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 and Intelligence-Exchange Alignment
----------	---

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.

Threat Intelligence is the core function that collects, catalogs, assesses and alerts on indicators and named threats from global sources, publishing each with a confidence and severity rating that is republished as the underlying evidence changes rather than set once and left.

A raw indicator on an external feed rarely says what it means. An address or a hash on its own does not say what campaign it belongs to, how fresh the sighting is, or whether another source corroborates it, so an analyst has to research each one by hand before deciding whether to act. Different feeds tag and rate the same indicator differently, a global surge in one attack type is not obviously relevant to a specific tenant until someone checks, and a severity assigned once is rarely revisited as the picture around it changes.

Threat Intelligence answers this with a five-stage curation mechanism. Indicators are tagged by type and category the moment they are collected from probes and third-party feeds, cataloged by feed, recency and geographic pattern, assessed against a defined severity taxonomy, and issued as an alert carrying impact, affected systems and recommendations. The rating attached to a published indicator or named threat updates continuously, republished as new corroboration or exploitability information arrives.

Global visibility only becomes actionable once it is checked against a specific tenant's own indicator history, so a worldwide pattern is surfaced as locally relevant rather than left as an undifferentiated feed dump.

THE FUNCTION IS PRECISELY SCOPED.

Threat Intelligence collects, curates and rates indicators and named threats. It does not author or validate detection logic, which is Detection Factory. It does not compute the risk score, which is the Risk Scoring Formula. And it does not decide or execute a response, which is SOAR and Playbooks. Each of those is a separate core function with its own brief.

Five stages

Collect, catalog, assess, alert, publish

Global and local

Worldwide patterns checked against tenant exposure

Never frozen

Confidence and severity are re-rated

SECTION 2

Why Unmanaged Feeds Fall Short

Threat intelligence consumption is underinvested relative to the rest of the security operation, and the consequences surface as missed context rather than as a single visible failure. Four failure modes explain why, and each is answered by a stage of the mechanism in section 3.

2.1 Indicators without context

A raw indicator on a feed does not say what campaign or vulnerability it belongs to, how fresh the sighting is, or what to do about it, so an analyst has to research each one by hand before deciding whether it matters.

2.2 Severity that is not consistently applied

Without a shared level taxonomy, one analyst's critical is another's medium, so the same class of threat is rated inconsistently across an alert history depending on who assessed it.

2.3 Global patterns invisible at the tenant level

A worldwide surge in a specific attack type is only actionable once a tenant can see whether their own indicators show the same pattern, and a raw feed dump does not distinguish between trending everywhere and specifically probing this environment.

2.4 Advisories that go stale

A severity assigned once, to an indicator or to a named vulnerability, is rarely revisited as exploitability and affected-system information change, so an old critical rating can outrank a genuinely current threat simply because nobody downgraded it.

2.5 The cost of the workaround

The usual response is an analyst cross-referencing several external feeds and vendor advisories by hand before triaging anything, a task that scales with the number of feeds subscribed to and does not get easier as feed volume grows. The underlying problem, indicators and advisories that are not curated, rated and kept current from one place, is untouched.

Industry research on security operations consistently finds analyst time lost to cross-referencing threat feeds by hand, and a meaningful share of indicators reaching a security operation uncorroborated by a second source, though the specific figures are not yet sourced for this brief.

SECTION 3

How Threat Intelligence Works

Threat Intelligence is one continuously operating mechanism native to the TDIR core, not a feed subscription a person monitors. Raw indicators and vulnerability disclosures enter at one end, and a curated, rated, continuously refreshed advisory leaves at the other.

Figure 1. The five-stage curation mechanism.

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
Collect	Indicators of attack are ingested in real time from global probe telemetry and third-party feeds, tagged by type and category at the point of intake.	Indicators without context
Catalog	Indicators are organized by feed source, recency and multi-source activity, and compared across local and global scope, so a worldwide pattern can be checked against a tenant's own indicator history.	Global patterns invisible at the tenant level
Assess	A severity is set against a defined level taxonomy. Where the threat is a named vulnerability, exploitability, classification and affected systems are attached alongside its severity scoring.	Severity that is not consistently applied
Alert	The assessed threat is issued as an advisory carrying its title, description, impact, affected systems and recommendations, and is tracked in the alert history.	Indicators without context
Publish	The curated indicator or named threat, with its rating, is published into the shared model, and is republished as new corroboration or exploitability information arrives.	Advisories that go stale

The stages run in sequence, but Publish is never final: a rating set at Assess is revisited the moment Collect or Catalog surfaces new evidence bearing on the same indicator or threat.

3.2 How a rating 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 tenant. 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
Indicator freshness	Published indicators reconfirmed within a defined recency window, as a share of the catalog.	Prioritized re-collection, or retirement of the stale indicator.
Multi-source corroboration	Published indicators confirmed by more than one feed, as a share of the catalog.	A confidence downgrade for single-source indicators, pending corroboration.
Severity consistency	Alerts whose assigned severity matches the defined criteria for that classification, as a share of alerts reviewed.	A review of the rating against the defined taxonomy.
Time to publish	Elapsed time from an indicator's first collection to being published with confidence and severity attached.	Investigation of the stage where the latency is introduced.
Local relevance rate	Published threats confirmed to touch a specific tenant's own indicator history, as a share of all published threats.	A review of the matching criteria against that tenant's indicator history.

A downgraded rating is the mechanism working, not a mistake to explain away. *An indicator whose confidence falls as corroboration fails to materialize, or a severity that drops as exploitability information is updated, is Publish doing exactly what it is built to do. Treating every downgrade as an error to investigate misreads a rating that is supposed to move.*

SECTION 4

Threat Intelligence in the Platform

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

What it draws on. Real-time probe telemetry and third-party threat-intelligence feeds as external sources, and the normalized entity model the Centric-AI Fabric carries, used to check whether a global pattern touches a specific tenant's own infrastructure. Each is an input to a specific stage rather than a general capability the function generates itself.

What it writes back. The curated indicator or named threat, with its confidence and severity rating, published into the shared model. The Detection Factory's source stage consumes it as raw material for new detection content, and the Risk Scoring Formula's threat-intelligence dimension consumes the rating directly. The relationship is one way: nothing returned to Threat Intelligence changes a rating already published, and a later re-rating comes from new evidence arriving at Collect or Catalog.

What it does not do. It does not author or validate detection logic, compute a risk score, or decide or execute a response. Threat Hunting may draw on the same curated indicators when forming a hypothesis to test, but Threat Intelligence does not itself investigate a tenant's environment. It holds no tenant telemetry beyond what is needed to check local relevance against the Centric-AI Fabric's model.

4.2 A worked example

The chain below is illustrative and technically representative rather than drawn from a named engagement. A vulnerability affecting a widely used VPN appliance is disclosed externally. Collect ingests the first exploitation indicators tied to it, tagged to the disclosure. Catalog links the pattern to the tenant's own indicator history and confirms probes from a known campaign region reaching the tenant's edge, not only a global trend. Assess assigns severity using the published CVSS metrics and the defined level taxonomy, noting the affected systems and versions. Alert issues the advisory with impact and recommendations. Publish sends the confidence-scored, named threat into the shared model, where the Detection Factory drafts a detection for the exploitation pattern and the Risk Scoring Formula's threat-intelligence dimension reflects the new campaign.

4.3 Delivery across many tenants

A threat curated once from global sources is checked against every tenant's own indicator history independently, so the same advisory can be locally relevant to one tenant and merely informational to another. What travels between tenants is the curated indicator and its global assessment; what never travels is one tenant's own indicator history or infrastructure exposure.

Because the function is switched on with the platform rather than deployed separately, time to value is a matter of feed coverage rather than an onboarding project. In the first week, global probe telemetry and the feeds already connected begin populating the catalog, and a tenant's own indicator history starts accumulating the moment the platform is live. Through the first month, local-relevance matching sharpens as that history deepens. At steady state, indicator freshness and multi-source corroboration are tracked as feed-health measures in their own right, reviewed whenever a source is added or retired.

SECTION 5

ATT&CK and Intelligence-Exchange Alignment

Threat Intelligence does not map detection content and does not compute technique relevance itself. Where an indicator or named threat is later translated into detection logic, the Detection Factory attaches the ATT&CK technique identifier at that point. This function's own contribution is the indicator, its confidence and its named-threat context, not the technique mapping.

FRAMEWORK OR STANDARD	ITS ROLE IN THREAT INTELLIGENCE
STIX and TAXII	Structured intelligence exchange for indicator and named-threat intake.
CVE and CVSS	Vulnerability identification and severity scoring for named threats in the vulnerability-oriented advisory catalog.
MITRE ATT&CK, Enterprise, Cloud and Identity	Carried forward once the Detection Factory translates a curated indicator into detection content. It is not mapped or scored by this function.
NIST Cybersecurity Framework 2.0	Outcome language used where platform reporting discusses threat-landscape posture to an executive or board audience.

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 Threat Intelligence compute the risk score or author detections?	No. Those are the Risk Scoring Formula and the Detection Factory. Threat Intelligence supplies the indicator, its confidence and the named-threat context those functions consume.
How is severity decided?	Against the defined levels the platform's Threat Level Manager maintains, with CVSS metrics attached where the threat is a named vulnerability.

QUESTION, AS A BUYER ASKS IT	ANSWER
What stops one feed's noisy indicator from being treated as high confidence?	Multi-source corroboration is tracked. An indicator confirmed by only one feed is rated accordingly.
Does a global threat automatically apply to every tenant?	No. Global visibility is checked against a tenant's own indicator history before anything is surfaced as locally relevant.
How does an outdated severity rating get corrected?	Publish is continuous. A rating updates as new corroboration or exploitability information arrives, rather than staying frozen at first assessment.
Can a tenant or provider add its own threat entries?	Yes. An authorized user can add a named threat or adjust a severity level, which then flows through the same Assess and Alert stages as an externally sourced one.

SECTION 7

The Benefits

Curated intelligence is where a feed subscription becomes an operational input, and the outcomes are countable: how much analyst time goes to judgment rather than cross-referencing, whether a rating can be trusted to be current, and whether a global pattern can be shown to touch this particular environment.

7.1 At a glance

BENEFIT	WHAT PRODUCES IT
An advisory arrives ready to act on rather than ready to research	Tagging at Collect and the impact, affected systems and recommendations attached at Alert.
The same class of threat is rated the same way every time	One severity taxonomy applied at Assess, rather than an individual analyst's judgment.

BENEFIT	WHAT PRODUCES IT
A rating can be trusted to be current	Publish is continuous, so confidence and severity move as corroboration and exploitability information arrive.
A worldwide pattern is separated from a local one	Catalog checks the global picture against each tenant's own indicator history before anything is surfaced.
Feed quality becomes visible rather than assumed	Indicator freshness and multi-source corroboration, defined in section 3.2 and reported per tenant.
One curation effort serves every tenant	A threat curated once from global sources is matched for local relevance independently, per tenant.

7.2 What changes, and for whom

Curated, continuously current intelligence changes what each role spends time on, and what each role can trust.

ROLE	WHAT CHANGES
SOC analyst	Advisories arrive with context, severity and recommendations already attached, rather than as a raw feed entry to research from scratch.
SOC manager	Indicator freshness and corroboration become visible, manageable numbers rather than a general sense of whether the feeds are trustworthy.
CISO and risk	Every severity rating traces to a defined, consistent taxonomy, which supports board and audit reporting on threat-landscape posture.
Service provider practice lead	A global threat curated once is checked for local relevance across every tenant independently, without a separate curation effort per customer.

7.3 What sets it apart

- **Tagged at the point of collection, not after.** Type and category are attached when an indicator first arrives, rather than reconstructed later by an analyst.
- **Global visibility, tenant-scoped relevance.** A worldwide pattern is surfaced as actionable only once it is checked against a tenant's own indicator history.

► **Continuously republished, never frozen.** Confidence and severity update as new corroboration or exploitability information arrives, which is what keeps an old critical rating from outranking a current threat.

THE VALUE IN ONE LINE.

An unmanaged feed hands a security operation a list of indicators and leaves the meaning to an analyst. A curated one hands over a rated, named, locally checked threat, and keeps the rating honest as the evidence moves.

SECTION 8

Summary

Threat Intelligence is how the platform turns raw external indicators and vulnerability disclosures into curated, rated, continuously current advisories. It is one core function among twelve, scoped to a five-stage curation mechanism that runs continuously inside the TDIR core rather than as a feed subscription someone monitors by hand.

- Collect, catalog, assess, alert and publish replace unmanaged, context-free feed consumption.
- Every indicator is tagged at collection, and every advisory is rated against one consistent severity taxonomy.
- Global patterns are checked against each tenant's own indicator history before being surfaced as locally relevant.
- Every published rating stays live, revisited as new evidence arrives rather than frozen at first assessment.
- The function curates and rates threat data, and leaves detection authoring, scoring and response 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.