

CORE FUNCTION BRIEF

Risk Scoring Formula

A core function of the ClearSkies iISOC platform

*One explainable value, correlated across seven dimensions,
continuously recalculated.*

Contents

1 Executive Summary

2 Why Static Severity Fails

A label fixed at the moment of detection

Severities that do not mean the same thing twice

Weak signals that never accumulate

Re-triage as the only way risk updates

The cost of the workaround

3 How the Risk Scoring Formula Works

The seven dimensions

How urgency is judged

Continuous re-scoring

4 The Risk Scoring Formula 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 Supported Frameworks

6 Business Impact and Delivery

What changes, and for whom

What sets it apart

Objection handling

7 Summary

SECTION 1

Executive Summary

IN ONE SENTENCE.

The Risk Scoring Formula is the core function that converts every alert into a single, explainable value by continuously correlating seven risk dimensions, so urgency reflects current, correlated risk rather than a severity label fixed at the moment of detection.

A conventional security operation asks one static question of every alert: is this high or low severity? That label is fixed at detection time, blind to context, and identical whether the affected asset is a test server or a domain controller. It does not move as an investigation reveals more, and it means something different in every tool that assigns it.

The Risk Scoring Formula replaces the label with a process. Behavioral, technique, asset, identity, intelligence, temporal and environmental signals enter as seven correlated dimensions. Each carries a dynamic weight, tenant-specific rather than fixed, and the platform recalculates the resulting value the moment any dimension changes. The value falls into one of four bands, Low, Medium, High and Critical, and each band changes what becomes possible next without the formula itself performing the next step.

Two consequences follow. An alert from an identity provider and an alert from an endpoint agent land on the same scale, because both are scored against the same seven dimensions and the same shared entity model. And a decision made a minute ago is revised the moment new evidence arrives, so a slow-burning pattern that stays below any single tool's threshold still climbs the scale as its signals accumulate.

The formula is native to the TDIR engine rather than a module added to it. Every alert the platform ingests, from every layer and every native add-on, passes through the same seven-dimension mechanism before anything downstream acts on it, so the score in a monthly report and the score on a single alert come from one calculation rather than from two tools reconciled by hand.

THE FUNCTION IS PRECISELY SCOPED.

The Risk Scoring Formula decides how urgent an alert is. It does not decide who handles it, which is Intelligent Alert Routing, and it does not decide or execute what happens next, which is SOAR and Playbooks. Each of those is a separate core function with its own brief.

7 dimensions

Correlated, dynamically weighted

1 explainable value

Continuously recalculated

One scaleConsistent across every layer

SECTION 2

Why Static Severity Fails

Prioritization is 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 dimension of the mechanism described in section 3.

2.1 A label fixed at the moment of detection

Severity is set once, when a rule fires, and is never revisited. It is identical whether the affected asset is a test server or a domain controller, and it does not change as an investigation reveals more or less about the underlying activity. The measurement that would reveal the drift, whether the original label still fits, belongs to nobody.

2.2 Severities that do not mean the same thing twice

A high severity in one tool reflects that tool's own internal scale, calibrated to its own detection logic. Across a SIEM, an identity provider and an endpoint agent, three unrelated high severities are not comparable, so an analyst prioritizing across consoles is comparing labels rather than risk.

2.3 Weak signals that never accumulate

A single low-risk event is flagged or ignored at the moment it occurs, and nothing in a static model tracks whether the same weak signal is recurring. A low-and-slow pattern that stays below any one alert's threshold can persist for weeks with no rising score to surface it. The analyst who absorbs the resulting noise has no way to see the pattern building underneath it.

2.4 Re-triage as the only way risk updates

When new evidence changes the picture, a fresh threat-intelligence match, a privilege escalation, or an asset reclassified as critical, a static severity does not move on its own. An analyst has to notice the change and manually re-triage, so most alerts are prioritized on the information available the moment they fired rather than on the fuller picture available minutes later. The organization stays exposed to a risk it has, in effect, already confirmed.

2.5 The cost of the workaround

The usual response is more headcount rather than a different mechanism: analysts added to hold several consoles open and reconcile unrelated severity scales by memory. Cost scales with alert volume and with the number of tools in the stack, while the underlying problem, a label that cannot update itself, is untouched

The workaround also fails quietly rather than visibly. A team that adds an analyst to cover the reconciliation work does not see the missed slow-burning pattern in a staffing report. It sees only the incident review after the fact, by which point the cost has been paid in dwell time rather than in headcount.

SECTION 3

How the Risk Scoring Formula Works

The Risk Scoring Formula is one continuously operating mechanism native to the TDIR engine, not a rule evaluated once. Seven dimensions feed one score, the score falls into one of four bands, and the whole calculation runs again whenever new evidence arrives. Figure 1 shows the model and the re-scoring loop.

Figure 1. The seven-dimension model and the continuous re-scoring loop. Re-scoring is what keeps the value current as evidence changes.

3.1 The seven dimensions

Each alert is evaluated across every dimension, never on one in isolation. The table names what each dimension measures, the platform source that supplies its signal, and the failure mode from section 2 that it answers.

DIMENSION	WHAT IT MEASURES, AND WHERE THE SIGNAL COMES FROM	FAILURE MODE IT ANSWERS
Behavioral confidence	Deviation from established baselines and anomalous behavior over time, such as abnormal login times, rare command execution or unusual access paths. Sourced from UEBA.	Weak signals that never accumulate
ATT&CK relevance	Alignment with known techniques and kill-chain progression, such as credential access chaining into lateral movement. Sourced from the technique identifier the Detection Factory attaches to the firing detection.	Severities that do not mean the same thing twice

DIMENSION	WHAT IT MEASURES, AND WHERE THE SIGNAL COMES FROM	FAILURE MODE IT ANSWERS
Asset criticality	Business and operational importance of the affected asset, such as domain controllers, cloud administrator accounts, operational technology or crown-jewel applications. Drawn from the shared entity model.	A label fixed at the moment of detection
Identity risk context	Privilege level, account type and behavior of the identity involved, including service accounts, dormant identities and multi-factor authentication status.	A label fixed at the moment of detection
Threat-intelligence confidence	Correlation with internal or external intelligence, including known malicious addresses and domains, active campaigns and industry-specific threats. Sourced from Threat Intelligence.	Re-triage as the only way risk updates
Temporal and pattern persistence	Whether activity is isolated or recurring, such as low-and-slow persistence or repeated failed actions over days or weeks.	Weak signals that never accumulate
Environmental context	Whether behavior occurs inside or outside expected operating conditions, such as maintenance windows, business hours or planned change events. It lowers the score during an expected anomaly as readily as it raises it.	Severities that do not mean the same thing twice

Risk increases as dimensions correlate rather than in isolation. A single low-risk dimension rarely escalates an alert on its own. Several reinforcing dimensions together drive it up the scale, which is how the formula surfaces a genuine slow-burning pattern while keeping routine noise quiet.

3.2 How urgency is judged

The score is judged against four bands, and each band changes what becomes possible downstream without the formula itself performing the next step.

BAND	WHAT CROSSES INTO THIS BAND	WHAT IT ENABLES DOWNSTREAM
Low	The score stays below the review threshold.	Intelligent Alert Routing keeps the alert aggregated rather than assigning it to an analyst queue.
Medium	The score crosses the review threshold.	Intelligent Alert Routing assigns the alert to an L1 analyst tier.
High	The score crosses the priority threshold.	Intelligent Alert Routing assigns with priority to an L2 or L3 analyst, and SOAR flags the case for expedited handling.

BAND	WHAT CROSSES INTO THIS BAND	WHAT IT ENABLES DOWNSTREAM
Critical	The score crosses the containment threshold.	The governance policy in SOAR evaluates the case for autonomous containment. Where authorized, Playbooks execute while Intelligent Alert Routing maintains oversight.

A low band is not zero risk. A Low band means the current evidence does not yet justify analyst attention, not that the event is dismissed. The event remains part of the correlated picture, and a recurring or escalating pattern raises the score as further signals arrive, which is the mechanism section 2.3 describes as missing from a static model.

3.3 Continuous re-scoring

Risk is never static, so the score is recalculated whenever the picture changes. The formula re-scores when new telemetry arrives, when additional ATT&CK techniques are detected, when identity privileges change, when asset importance is updated, and when threat intelligence is refreshed. Three things follow from that.

- **Slow attacks escalate early.** Low-and-slow activity that would evade a static severity climbs the scale as its signals accumulate.
- **False positives de-escalate.** Alerts that prove benign fall back down the scale without consuming analyst time.
- **Urgency adapts without manual re-triage.** The band follows the risk automatically as conditions change, rather than waiting for someone to notice.

SECTION 4

The Risk Scoring Formula in the Platform

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

What it draws on. The technique identifier and precision history the Detection Factory attaches to a firing detection; behavioral baselines and anomaly measurements from UEBA; indicator confidence and campaign relevance from Threat Intelligence; and the asset, identity and environmental context the shared entity model carries. Each is named as an input to one of the seven dimensions rather than as a general capability the formula generates itself.

What it writes back. A score and an urgency band, published into the shared model. Intelligent Alert Routing, SOAR and Playbooks consume that output. The relationship is one way: nothing is returned to the Risk Scoring Formula that changes how the current score was calculated, and a later re-score is

triggered by new evidence arriving at one of the seven dimensions rather than by a signal a downstream function sends for the purpose.

What it does not do. It does not choose an analyst, a queue or a response action, and it holds no telemetry of its own. Its contribution is not a new signal but a single, comparable answer to how urgent an existing signal is.

Beyond the bands themselves, the same value feeds analyst assignment by skill level and workload, escalation that respects the service levels in force, and immediate handling for high-confidence threats. In every case the formula supplies the value and a sibling core function performs the action.

4.2 A worked example

FROM A QUIET FIRST SIGNAL TO A CRITICAL SCORE.

A service account authenticates from an unfamiliar location outside business hours. On its own the event scores modestly: identity risk context and environmental context contribute, while ATT&CK relevance and behavioral confidence stay low because no technique has matched and the account's baseline tolerates occasional off-hours activity. Minutes later, Detection Factory content flags a technique match for credential use consistent with a known campaign, and UEBA reports the same account performing a rare directory-enumeration command. Behavioral confidence and ATT&CK relevance both rise, and because the signals reinforce rather than stand alone, the score crosses from Medium to High. Threat Intelligence then confirms the source network range is associated with active campaign infrastructure, and the score crosses into Critical. The single off-hours login, considered alone twenty minutes earlier, would not have justified any of this.

4.3 Delivery across many tenants

Weighting is configured per tenant, industry and asset group, so the same seven-dimension mechanism adapts to each client's environment without a bespoke rebuild. What travels between tenants is the weighting configuration and the scoring logic. What never travels is the telemetry that produced any tenant's score. A provider can therefore explain the reasoning behind an urgency decision to any client, from the same mechanism, without a manual audit or a separate scoring engine per tenant.

Every change to a tenant's weighting configuration is logged with its date, its author and the reason recorded against it, so a shift in how an environment is scored is traceable rather than silent. A weighting change is reviewed before it takes effect, and nothing adjusts the dynamic weights in production without passing through that review.

SECTION 5

ATT&CK Alignment and Supported Frameworks

ATT&CK relevance is one of seven correlated dimensions, not the sole basis for a score. The technique identifier a given alert carries is supplied by the Detection Factory at the point of detection. The Risk Scoring Formula consumes that identifier, together with kill-chain progression across chained techniques, as one input among seven. A technique is cited by identifier and name on first use, for example T1078 Valid Accounts, and by identifier alone thereafter.

Technique relevance is scored across the Enterprise, Cloud and Identity matrices, matching the matrices the Detection Factory maps detection content against. Where an alert chains multiple techniques, later positions in the chain, closer to impact than to initial access, weigh more heavily on this dimension, which is what allows a discovery step and a collection step to score differently even when both map to a technique in the same matrix

Mapping is not coverage. A technique listed in the coverage matrix is claimed. A technique exercised in a purple-team run or against MITRE Engenuity ATT&CK Evaluations is demonstrated. The distinction belongs to the Detection Factory, which owns the matrix. The Risk Scoring Formula consumes the identifier rather than validating it.

FRAMEWORK OR STANDARD	ROLE IN THE RISK SCORING FORMULA
MITRE ATT&CK, Enterprise, Cloud and Identity	Supplies the technique identifier and kill-chain position that inform the ATT&CK relevance dimension
NIST Cybersecurity Framework 2.0	Outcome language used where platform reporting discusses aggregate risk posture to an executive or board audience.

Alignment to OCSF and OpenTelemetry, and the syslog, CEF, LEEF and REST API ingestion formats, are properties of the Collection Layer and are described in the *ClearSkies iISOC platform brief* rather than restated here.

SECTION 6

Business Impact and Delivery

The Risk Scoring Formula reports the measures below natively and per tenant. Definitions are published because the same metric name is measured differently by different vendors. Target values are agreed per engagement rather than asserted here.

MEASURE	DEFINITION USED	REPORTED AS
Band accuracy	Critical-band and High-band alerts confirmed as true positives on investigation, as a share of all alerts in that band.	Trended per tenant.
Re-scoring latency	Elapsed time from new evidence arriving at any dimension to the score reflecting it.	Trended, with outliers flagged.
De-escalation rate	Alerts that fall back to the Low band as further evidence proves them benign, as a share of all scored alerts.	Reported as analyst time protected from unnecessary triage.
Escalation lead time	For a confirmed slow-burning pattern, the interval between the first weak signal and the score crossing into High or Critical.	Reported per confirmed case.
Dimension coverage	Scored alerts with all seven dimensions populated with a non-default value, as a share of all scored alerts.	Per tenant, with the weakest dimension named.

6.1 What changes, and for whom

ROLE	WHAT CHANGES
SOC analyst	Alerts arrive with a current, correlated urgency rather than a fixed label, so triage time goes to genuine risk rather than to reconciling severities across consoles.
SOC manager	Prioritization quality becomes a reportable, explainable number rather than a judgment call that varies by analyst and by shift.
CISO and risk	Every urgency decision carries a defensible rationale, current at the moment it is read rather than at the moment the alert first fired.
MSSP practice lead	One scoring mechanism applies the same explainable logic across every tenant, and weighting adapts per client without a bespoke engine.

For the organization defending its own operations, effort concentrates where business impact is greatest, because the score reflects asset criticality and identity context rather than the opinion encoded in a rule. For the service provider, the same mechanism removes manual triage as the bottleneck and lets a lean team deliver predictable outcomes across many tenants.

6.2 What sets it apart

- **Seven correlated dimensions, one explainable value.** No single weak signal escalates an alert; reinforcing signals across dimensions do.
- **Dynamic rather than fixed weighting.** Weight ranges adjust per industry, tenant and asset group, and recalculate as conditions change.
- **Continuous re-scoring.** A score changes the moment new telemetry, intelligence, identity or asset context arrives, rather than waiting for manual re-triage.
- **One scale across every layer.** An identity event and an endpoint event are ranked on the same scale because both are scored against the same seven dimensions.

6.3 Objection handling

QUESTION, AS A BUYER ASKS IT	ANSWER
Does the Risk Scoring Formula decide which analyst handles an alert?	No. It computes the score and the urgency band. Intelligent Alert Routing reads that output to decide the queue, the analyst tier and the assignment by skill and workload.
Does the score itself trigger autonomous containment?	No. A Critical band makes autonomous containment eligible. The governance policy in SOAR decides whether it is authorized for that class of case, and Playbooks execute the action once authorized.
How is the score different from a single-vendor severity rating?	A single-vendor severity is fixed at the point a rule is defined and reflects one dimension. The Risk Scoring Formula correlates seven dimensions, weights them dynamically per tenant and asset group, and recalculates continuously as new evidence arrives.
Can the weighting be tuned for a specific industry or client?	Yes. Weight ranges are adjustable per industry, tenant and asset group and are recalculated as conditions change, so the same mechanism adapts without a bespoke rebuild. Every change is logged and reviewed before it takes effect.
What happens when new intelligence arrives after an alert has already been scored?	The score recalculates immediately. A refreshed threat-intelligence match, a privilege change or an asset reclassification all trigger re-scoring without waiting for manual re-triage.

**QUESTION, AS A
BUYER ASKS IT****ANSWER**

Does it receive anything back from alert routing, SOAR or playbooks?

No. The relationship is one way. The Risk Scoring Formula publishes the score and the band, and those functions consume them. A later re-score is triggered by new evidence, never by a return signal.

SECTION 7

Summary

The Risk Scoring Formula is how the platform decides how urgent an alert is. It is one core function among twelve, scoped to computing and continuously recalculating a single, explainable value by correlating seven risk dimensions, and it runs inside the TDIR engine rather than as a judgment made once and left to age.

- ▶ Seven correlated dimensions, sourced from the Detection Factory, UEBA, Threat Intelligence and the shared entity model, replace a severity label fixed at detection time.
- ▶ Weighting is dynamic and tenant-configurable, and every change to it is logged and reviewed before it takes effect.
- ▶ The score recalculates continuously as new telemetry, intelligence, identity or asset context arrives.
- ▶ Every alert, whatever layer it came from, is expressed on the same scale and in the same four bands.
- ▶ The function decides urgency, and leaves ownership and response to the core functions built for those decisions.

External research is cited inline where it is used, and this brief carries no separate source register. 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.