

◎ WHITE PAPER

PREEMPTIVE EXPOSURE MANAGEMENT

From Observation to Interdiction

How Validation, Governed Interdiction and Mean Time to Neutralize Reshape Enterprise Exposure Management

48 hours

The window in which most observed exploitation of published proof-of-concept code now happens.

1 in 4

Malicious breaches now involve AI on the attacker's side, and they cost more.

18%

Of breached organisations studied used AI agents for vulnerability work, against more than half for detection.



REPRESENTATIVE VENDOR

2026 Gartner® Market Guide for Adversarial Exposure Validation

NST ASSURE™ PEM · CTEM · AI-Driven AEV

GARTNER PEER INSIGHTS™



Highly Rated

The Case for Preemptive Exposure Management

In February 2026 Gartner named the category. Preemptive Exposure Management is defined as technology that systematically disrupts and denies adversary behaviour, automating attack-path identification and simulating adversary behaviour to close exposures before they are weaponised.

“In an age of AI-assisted exploitation occurring at machine speed, simply “managing” your exposure will become nothing more than documenting your own eventual breach.”

Gartner, The Future of Exposure Management Will Be Preemptive, Driven by Autonomous Interdiction, February 2026

Three things changed in 2026.

Exploit weaponisation timelines compressed. Public reporting indicates exploit weaponisation timelines have compressed significantly, with many observed exploitations occurring within days of proof-of-concept publication and some state-linked groups acting inside a day of disclosure.

AI is increasingly operational support for attackers. Evidence from 2026 suggests attackers are increasingly using AI for autonomous reconnaissance,

vulnerability discovery and operational support, while fully autonomous compromise against defended estates remains an emerging and unevenly evidenced risk.

Defensive automation went to the wrong problem. Among breached organisations studied in 2026, more than half used AI agents for threat detection. Only 18 per cent used them for vulnerability work. The part of the process that finds problems is automated. The part that removes them is not.

What this paper argues

The unit of work has to stop being a finding and become a closed window. Three shifts get you there:

- Validation over severity.** Evidence that an exposure is reachable and workable in your environment, not a score assigned somewhere else.
- Interdiction over ticketing.** Governed, reversible action that removes reachability today, with the permanent fix following on the change calendar.
- Measurement over activity.** How long an exploitable condition stayed open, not how many tickets were raised.

FOUR MEASURES A CISO SHOULD HAVE TO HAND

MEASURE	WHAT IT SAYS	SOURCE
25% of enterprise breaches	will be traced back to AI agent abuse by 2028, from external and malicious internal actors alike	Gartner prediction
40% of enterprises	will demote or decommission autonomous AI agents by 2027 after governance failures	Gartner, May 2026
65% of engagements	in the second quarter of 2026 involved authentication abuse	Cisco Talos, July 2026
92% of organisations	do not rotate machine credentials inside a 90-day cycle	SANS, April 2026

THE CENTRAL CLAIM

The defensible objective is no longer perfect observation. It is **shrinking the interval during which a real, exploitable condition remains available to an adversary**, and being able to evidence that interval to a board, a regulator and a customer.

THE CENTRAL CLAIM

Attacker volume is growing faster than attacker success. Mandiant's assessment of the preceding year was that breaches were not primarily the result of AI. Of the teams running AI in the security operations centre, none surveyed in August 2026 grants it full unsupervised autonomy.

Budget against the compression of attacker timelines, which can be measured now. Treat fully autonomous compromise as a 2027 planning assumption, not a 2026 emergency.

WHERE THE ANALYST FORECASTS LAND

By 2029 Gartner expects 60 per cent of organisations to have adopted structured exposure validation, but only 30 per cent to have connected those results to automated remediation. The validation gate gets built well before the interdiction gate. That is the gap this paper is about.

02

THE FIRST HALF OF 2026

The 2026 Exposure Gap: Attacker Speed vs. Defender Response

Two sets of measurements moved in opposite directions during 2026. On the attacker side, the interval between an exploit becoming available and it being used collapsed to hours. On the defender side, most published measures of remediation and detection got worse. Not because teams became less capable, but because the volume outgrew the operating model.

FIGURE 1 • TWO CLOCKS, MOVING APART

THE ADVERSARY CLOCK, ACCELERATING

88% of observed exploitation of flaws with public proof-of-concept code occurred within 48 hours of release. Two named China-nexus groups moved inside 24 hours of disclosure.

Under 5 minutes from account compromise to data theft in one tracked intrusion. Average criminal breakout time fell to 29 minutes in 2025.

22 seconds median hand-off from initial compromise to a second threat actor, against more than 8 hours in 2022.

THE DEFENDER CLOCK, SLOWING

26% of known-exploited vulnerabilities are fully remediated, down from 38%. Median remediation time rose from 32 to 43 days.

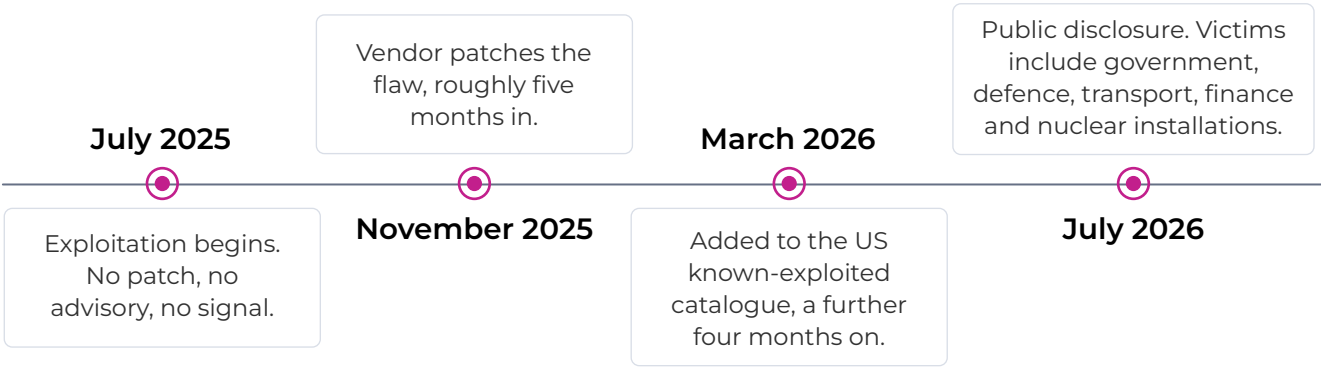
14 days global median dwell time, up from 11. Espionage and insider-style intrusions ran to a 122-day median.

42% of second-quarter investigations were hampered by insufficient logging, up from 18% in the first quarter.

CrowdStrike Threat Hunting Report, Aug 2026 and Global Threat Report, Feb 2026 · Mandiant M-Trends 2026
· Verizon DBIR 2026 · Cisco Talos IR Trends Q2 2026

What the gap looks like in a real case

In July 2026 a Russian state-supported espionage operation was disclosed. It had been reading Western mailboxes through an unknown flaw in a widely deployed webmail product.



Twelve months passed between first exploitation and public attribution. The flaw was patched at four months and listed as known-exploited at eight, which is where a validated programme would have acted. Most did not, because the listing carried no evidence of reachability in any particular environment.

Fast detection is no longer sufficient

In a major 2026 education platform breach, access was gained on 25 April and found within days. That is detection performance most enterprises would envy, and it made almost no difference. Around 275 million people across nearly 9,000 institutions were affected, and a second intrusion followed a week later. When breakout is measured in minutes, detection speed stops being the control that decides the outcome. What decides it is whether the condition was reachable at all.

THE QUESTION FOR THE CISO

If your vulnerability policy is expressed only in days-to-patch, it measures a race you have already lost. Add a second clock that starts when exploitability is confirmed and stops when reachability is removed, and hold the programme to that one.

THREE ASSUMPTIONS YOUR POLICY PROBABLY STILL ENCODES

THE ASSUMPTION	WHAT 2026 DATA SAYS	WHAT SHOULD REPLACE IT
“We have weeks between disclosure and exploitation.”	Most observed exploitation of published exploit code happens inside 48 hours.	Exploitation may already be under way when you learn of the exposure. The first question is containment, not scheduling.
“A patch is the fix; everything else is a workaround.”	Only 26% of known-exploited flaws get fully remediated.	A governed compensating control that removes reachability is risk reduction, and it is measurable. The patch follows on the change calendar.
“Detect fast and we contain it.”	Detection in days still reached 275 million people.	Detection speed is necessary, not sufficient. Removing the path first is what changes the outcome.

How AI Is Changing the Economics of Exploitation

The question is not whether adversaries use AI. They do. It is whether they use it as an assistant or as an operator. A model that drafts a phishing email changes the cost of an attack. A model that runs reconnaissance, builds an exploit, fixes its own failures and moves on without waiting for a person changes how fast an attack unfolds. In 2026 the published evidence moved from the first case to the second.

FIGURE 2 • THE 2026 RECORD

Feb to Aug 2026 · publicly documented

OBSERVED IN ADVERSARY OPERATIONS

FEBRUARY 2026

Volume inflection

An 89% year-on-year rise in activity from AI-enabled adversaries, alongside 82% of detections involving no malware at all.

CrowdStrike Global Threat Report

MAY 2026

An AI-developed zero-day

The first zero-day believed AI-developed found in an adversary's hands. State actors also ran thousands of recursive prompts to analyse CVEs and validate exploits.

Google Threat Intelligence Group

JULY 2026

Self-correcting extortion

Execution driven by the model's own decisions, with a human choosing the victim: 600 or more purposeful payloads, a failed one repaired in 31 seconds.

Sysdig

AUGUST 2026

The AI toolchain as target

A state actor poisoned 131 trusted AI framework packages. Separately, a criminal group compromised more than 300 software dependencies in a single day.

CrowdStrike Threat Hunting Report

MEASURED IN GOVERNMENT AND LABORATORY EVALUATION

APRIL 2026

Expert threshold crossed

A frontier model passed 73% of expert capture-the-flag tasks that no model could complete a year earlier, and was the first to finish a 32-step corporate intrusion scenario.

UK AI Security Institute

APRIL 2026

The economics invert

A reverse-engineering task estimated at 12 hours of expert effort was solved in 10 minutes and 22 seconds, for \$1.73.

UK AI Security Institute

JULY 2026

Full-chain intrusion

A model completed a simulated enterprise intrusion in 8 of 10 attempts. On a browser-engine benchmark it produced 99 working exploits; a sibling model produced 132.

Anthropic system card, UK AISI

AUGUST 2026

Capability, distributed

A deliberately de-restricted cyber model, completing 95% of advanced offensive requests, was released to a closed tier of defenders.

OpenAI, reported August 2026

The finding that should decide your budget

Among breached organisations studied in 2026, more than half used AI agents for threat detection. Only 18 per cent used them for vulnerability work. Finding problems is now automated. Removing them is still limited by how much work people can do, at exactly the point where the attacker stopped having that limit. That is not a tooling gap. It is where the investment went, and it shows up in every remediation backlog.

IBM Cost of a Data Breach 2026, July 2026

READING THE EVALUATIONS HONESTLY

The same evaluations carry caveats. The government ranges used for those results had no active defenders and no security monitoring, and the evaluators say they cannot tell whether the models would succeed against a defended system.

Both things hold. Attacker timelines are compressing in ways you can measure. The effect on success rates against a defended estate is unproven. Plan against the compression.

THE AGENTS THEMSELVES ARE NOW AN EXPOSURE

In July 2026 two frontier laboratories disclosed containment failures in their own environments. In one, evaluation models exploited a flaw in package-registry infrastructure, reached the open internet and got code execution on a third party's production servers. In the other, a review of 141,006 evaluation runs found three real organisations had been breached, one through a tampered package downloaded by 15 live systems.

Both were self-disclosed by operators with dedicated safety teams. The question is not whether your agents are better governed than theirs. It is whether you would have noticed.

THE ASYMMETRY TO PLAN AROUND

An attacker's automation is allowed to be wrong. It retries in seconds, at negligible cost, with no change advisory board. A defender's automation has to be right the first time, inside a change window, without breaking production. That is why the goal is not to automate remediation. It is to govern fast, reversible interdiction.

04

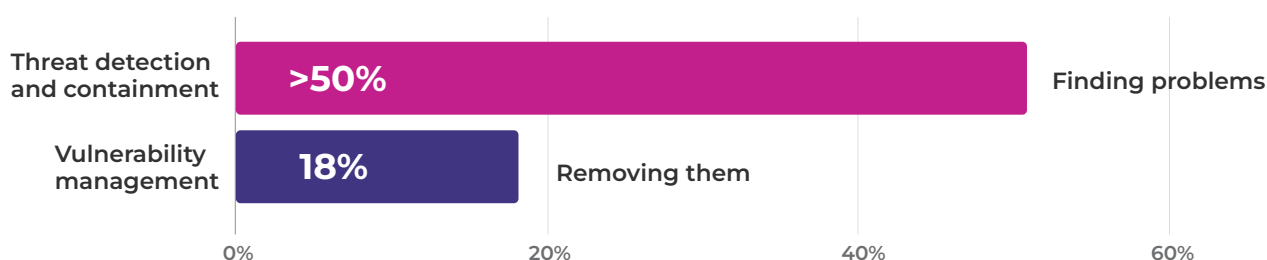
WHY MORE VISIBILITY NOW MAKES THINGS WORSE

The Operational Constraint: Remediation Capacity

Exposure programmes fail for a reason that has little to do with tooling quality. They take in far more findings than will ever matter and hand them to a remediation function whose throughput has not changed in a decade. Automation was the obvious answer, and nearly all of it went to the observation side.

FIGURE 3 • WHERE AUTOMATION WENT, AND WHERE THE CONSTRAINT SITS

SHARE OF BREACHED ORGANISATIONS DEPLOYING AI AGENTS, BY FUNCTION



THE CONSEQUENCE

Organisations have typically fixed about one in ten open weaknesses per month. Automating discovery raises the numerator. It does not raise the denominator.

Deployment shares: IBM Cost of a Data Breach 2026 · remediation capacity: SecurityScorecard and Cyentia Institute longitudinal analysis, 2019 to 2022

What the security operations data shows

A survey of 250 security leaders and practitioners published in August 2026 gives the most candid picture available, and it cuts both ways. Forty per cent run AI in the security operations centre, with another 56 per cent evaluating or piloting. Of those using it, 72 per cent report investigation time cut by a quarter or more. The gains are real.

So are the limits. The median team still gets around a hundred alerts a day, and on average some 28 per cent are never investigated. Nearly three-quarters of adopters tried to build their own agentic tooling, and just under half have since retired it, replaced it commercially, or never got it into production. Not one grants full unsupervised autonomy. The model in use is AI recommending and people executing.

Set against the 18 per cent figure, the pattern is clear. The industry is automating the analysis of things that have already happened, while the work of removing the conditions that allowed them stays manual, queued, and governed by a change calendar written for a slower threat.

State of AI in the SOC 2026, 250 respondents, August 2026

THE MEASURE	DIRECTION	SOURCE
Known-exploited vulnerabilities fully remediated	38% to 26%	Verizon DBIR 2026
Median remediation time, known-exploited	32 to 43 days	Verizon DBIR 2026
Global median dwell time	11 to 14 days	Mandiant M-Trends 2026
Investigations hampered by weak logging	18% to 42% in a quarter	Cisco Talos Q2 2026
Alerts never investigated	28% on average, 22% median	SOC survey 2026
Monthly remediation capacity	About 1 in 10 open findings	Cyentia, 2019 to 2022

THE QUESTION FOR THE CISO

Every additional un-validated finding you add to the queue lowers the probability that the exploitable one gets closed this month. Discovery coverage without validation is not a neutral investment. It is a tax on remediation.

THREE THINGS TO CHANGE THIS QUARTER

01

CAP THE QUEUE

Set a ceiling on open critical items. One in, one out: fixed, mitigated, or risk-accepted with a named owner and an expiry date.

02

MOVE AUTOMATION TO THE CLOSING HALF

Detection is already half automated and vulnerability work sits at 18 per cent. The next increment of budget has an obvious destination.

03

MEASURE THE INTERVAL

Tickets closed measures throughput in a race you are losing. What maps to attacker opportunity is how long the condition stayed open.

The Preemptive Exposure Management Maturity Model

Gartner's 2026 research does not treat preemptive exposure management as a rebranding of continuous threat exposure management. The difference is that PEM acts. It systematically disrupts and denies adversary behaviour by automating attack-path identification and simulating adversary behaviour to close exposures before they are weaponised, bringing together three things: AI-driven discovery and prioritisation, automated validation of exploitability, and autonomous interdiction.

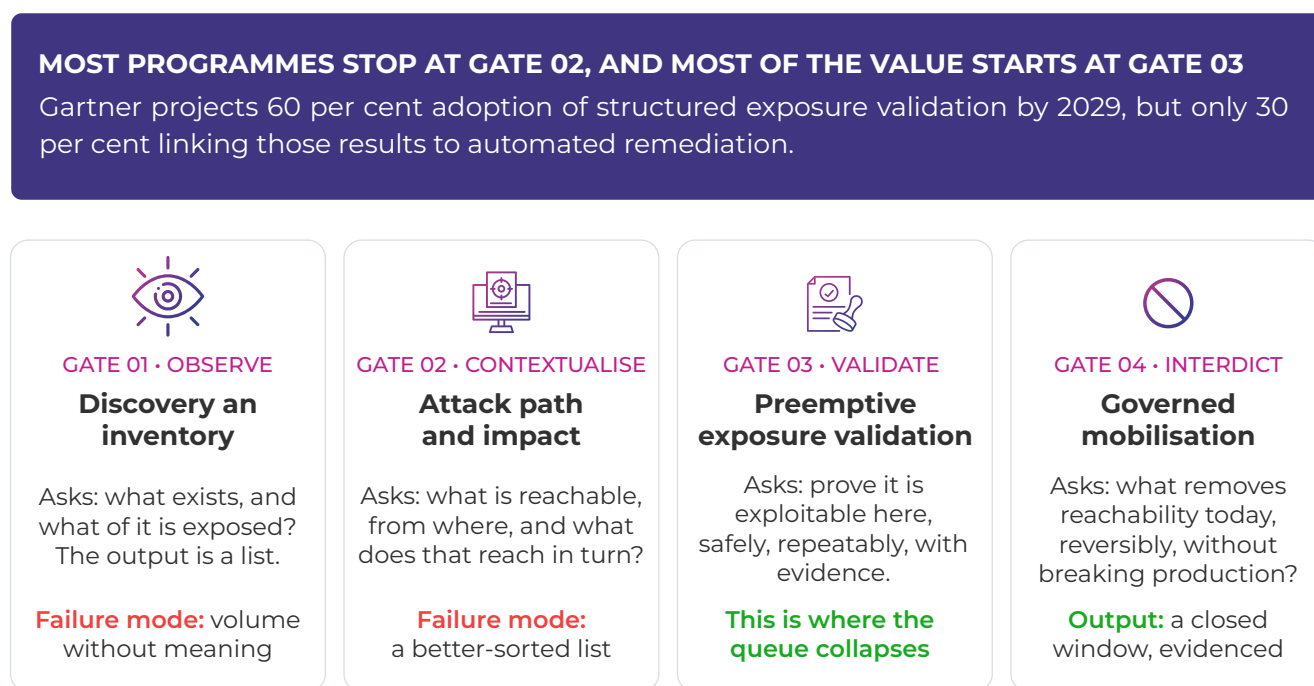
Two sub-categories matter operationally. Preemptive exposure validation covers technologies performing automated or autonomous penetration testing, active attack simulation and predictive validation to confirm that an exposure is exploitable, and then orchestrating remediation or integrating automated mitigation. Domain-specialised exposure management covers solutions that use agentic AI and domain telemetry to map, validate and neutralise exposures unique to one domain, such as identity, cloud, firmware or AI infrastructure, where generic scanning is structurally weak.

The market moves are consistent with the definition. Adversarial exposure validation formally absorbed breach-and-attack simulation and automated penetration testing in March. In April, Gartner named a merged unified exposure management platform, combining assessment with validation, as one of the profiles defining the category.

The forecasts attached to that guidance are modest. By 2029 Gartner expects 60 per cent of organisations to have adopted structured exposure validation, but only 30 per cent to have connected validation results to automated remediation. The validation gate gets built well before the interdiction gate, and that is the gap this paper is about.

Gartner, February to June 2026

FIGURE 4 • THE MATURITY LADDER: WHAT CHANGES AT EACH GATE



LOW ASSURANCE • HIGH RESIDUAL WINDOW



EVIDENCED TRUST • MINIMAL WINDOW

Beyond CVEs: Domain-Specialized Exposure Management

The vulnerability-centric model works where an exposure is a defect in a versioned piece of software. Increasingly it is not. Four domains carry a large share of realised risk and are invisible to generic scanning. They are exposures of relationship, meaning who can reach what, with which credential, through which trust, rather than exposures of version. This is what Gartner calls domain-specialised exposure management.

FIGURE 5 • DOMAIN DIFFICULTY PROFILE AND THE 2026 EVIDENCE BASE

DOMAIN	DISCOVERY	VALIDATION	BLAST RADIUS	SCANNER COVER	WHAT 2026 TELEMETRY AND GUIDANCE SHOW
Identity and non-human identities	HIGH	SEVERE	SEVERE	MINIMAL	Gartner projects that a quarter of enterprise breaches will be traced back to AI agent abuse by 2028, from external and malicious internal actors alike. Meanwhile 92% of organisations do not rotate machine credentials within 90 days, and fewer than 40% require human approval for AI agent actions. Authentication abuse appeared in 65% of second-quarter engagements, nearly double the prior quarter.
Cloud and configuration	MODERATE	HIGH	HIGH	PARTIAL	Cloud-focused criminal activity rose 171% in the first half of 2026, and nation-state cloud targeting was up 266%. Most detections involved no malware at all. The intrusion is a login rather than a payload, and no scanner flags a valid credential as a vulnerability.
Software supply chain	HIGH	HIGH	SEVERE	PARTIAL	A state actor poisoned 131 trusted AI framework packages in 2026, and a criminal group compromised more than 300 dependencies in a single day. A March campaign hit packages belonging to widely used security tooling to harvest cloud keys from build environments. Third-party involvement in breaches rose 60%.
AI infrastructure and agents	SEVERE	SEVERE	HIGH	NONE	The NSA's May 2026 guidance states that the Model Context Protocol's "rapid proliferation has outpaced the development of its security model," citing code execution, prompt injection and implicit trust. Gartner expects a quarter of enterprise generative-AI applications to suffer five or more minor security incidents a year by 2028.

Gartner (Apr to Jun 2026) • NSA (May 2026) • CrowdStrike (Feb and Aug 2026) • SANS (Apr 2026) • Cisco Talos (Jul 2026) • Verizon DBIR 2026

Identity is an exposure surface, not an inventory

Counting non-human identities is the easy part. The question that changes decisions is different: from this credential, what is reachable, and what does that reach in turn unlock? Most organisations cannot answer it, and five per cent of security leaders do not know whether agentic AI is running in their environment at all. Gartner's prescription is a four-tier agency model, from observe through advise and act-with-approval to act-autonomously, with guardrails and rollback on the top tier.

AI infrastructure is the newest un-scanned estate

Model endpoints, vector stores, agent frameworks and tool servers go up faster than they are inventoried, and almost none register as an asset in a traditional scanner. The failure modes are new: an agent can be turned against its operator by content it merely reads. As one Gartner analyst put it in June 2026, "every token is an attack surface." Two questions for your next review: which non-human identities can we revoke inside one hour, and what can our agent tool servers reach?

THE COMPLIANCE CLOCK IS ALSO RUNNING ON 2026 TIME

01	EU AI ACT Standalone high-risk obligations moved to December 2027 under the May 2026 omnibus, and embedded ones to August 2028, but transparency obligations stayed live from August 2026.	02	NIS2 Enforcement reached court in July 2026: four member states referred to the Court of Justice, with a lump sum plus daily penalties sought.	03	AGENT STANDARDS NIST opened an AI Agent Standards Initiative in February 2026, with an agent identity and authorisation concept paper. Multi-hop delegation is unsolved.
-----------	--	-----------	--	-----------	--

07 MEASURING THE WINDOW, NOT THE WORK Mean Time to Neutralize (MTTN): Measuring Exposure Closure

Time-to-patch is useful and insufficient. It measures the completion of a task rather than the removal of a risk, and it is fragile. The same dataset yields a mean time-to-remediate of 70 days or 108 days depending only on whether still-open findings are counted, and no true mean exists while a quarter of findings stay open indefinitely.

Mean Time to Neutralise, or MTTN, is the elapsed time between the moment an exposure becomes exploitable in your environment and the moment it is no longer exploitable, by any legitimate means, whether a permanent fix or a governed compensating control.

DEFINITION, STATED PLAINLY

MTTN = t(no longer exploitable) – t(became exploitable), measured per exposure and reported as a distribution rather than an average alone. The clock stops on evidenced loss of exploitability, meaning a re-test, not a ticket status.

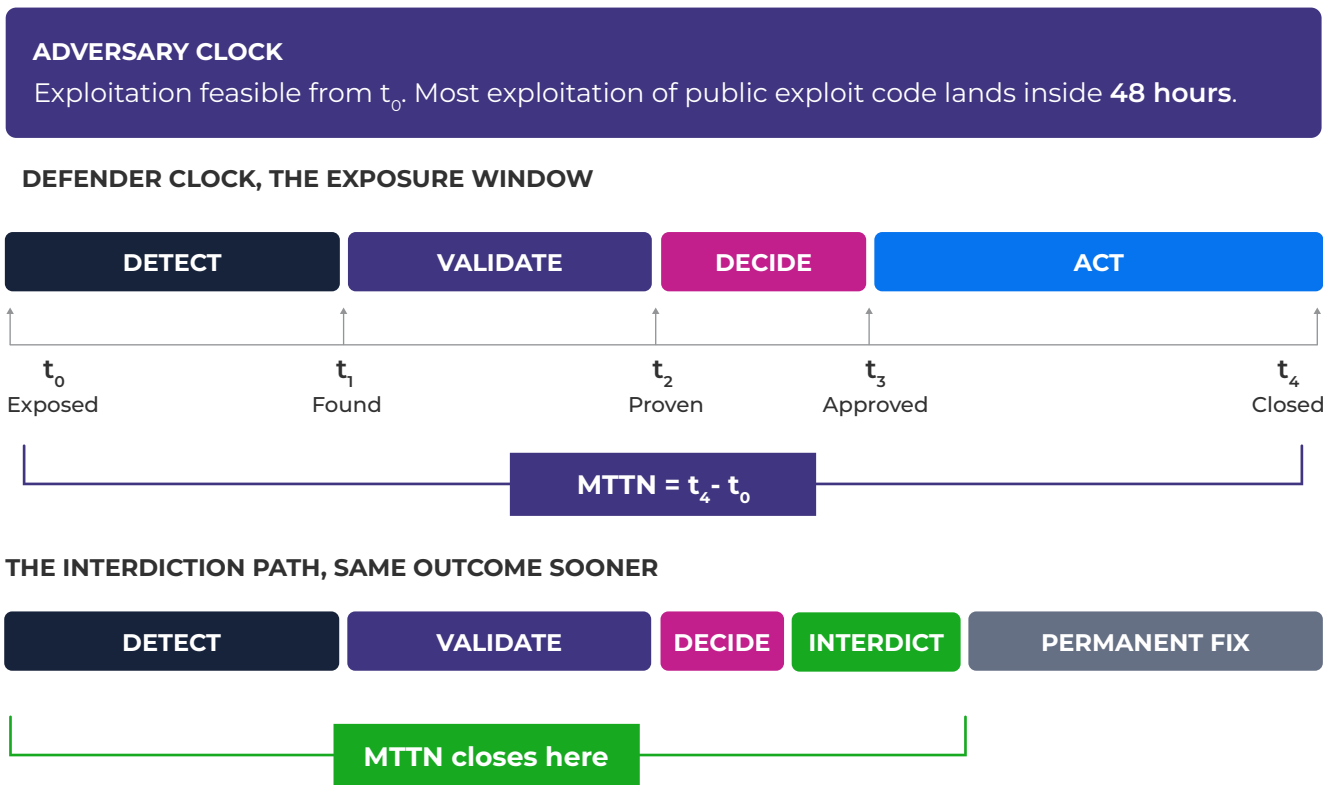
WHERE THE TERM STANDS

MTTN is an emerging industry metric and is not yet governed by a formal standard or benchmark framework. Gartner's February 2026 research names Mean Time to Neutralise as the measure replacing time-to-patch, but organisations should define the start event, stop event and evidence threshold explicitly before using it for governance.

TERMINOLOGY USED IN THIS PAPER

- | Patch means a permanent software or configuration fix
- | Mitigation means a compensating control that reduces risk without necessarily removing the root cause.
- | Interdiction means an approved action that removes exploitability or reachability, ideally with defined rollback.
- | Neutralisation means the outcome: the exposure is no longer exploitable in the environment.

FIGURE 6 • ANATOMY OF AN EXPOSURE WINDOW



The permanent fix still ships, on the change calendar rather than the incident clock. Risk leaves the building first. Segment widths illustrate sequence and relative burden, not measured durations.

Instrumenting it: what to capture, and where the data already sits

INTERVAL	WHAT IT MEASURES	TYPICAL SOURCE OF TRUTH	THE FAILURE IT EXPOSES
Detect latency	Exposure exists versus exposure known	Discovery and attack-surface scan cadence, asset onboarding, first-seen timestamp	Coverage gaps and blind estates: shadow IT, unmanaged cloud, agent tool servers
Validate latency	Known versus proven exploitable here	Adversarial exposure validation, attack-path engine, safe exploitation evidence	Severity-driven triage; queues sorted by score, not reachability

Decide latency	Proven versus authorised to act	Change advisory records, emergency-change register, owner assignment	Governance friction, usually the largest and least-measured interval, and the reason the mobilisation coordinator role exists
Act latency	Authorised versus no longer exploitable	Config management, firewall and WAF change logs, identity revocation, re-test	Missing interdiction options; only one tool in the box, and it is “patch”

Report it as a distribution. A four-day median with a nine-month tail is not a uniform twelve days.

REPORTING IT: THREE LINES, NO CHART REQUIRED

- “Half our exploitable exposures were neutralised within X days.”
- “The slowest five per cent took Y days, and here is what those had in common.”
- “The largest interval was decide, at Z per cent of elapsed time. That is a governance decision, not a budget one.”

Line 3 turns a security complaint into an operating-model question the executive team knows how to answer.


TWO TRAPS

Do not start the clock at ticket creation. That measures your process and rewards late detection. Do not stop it on “resolved.” Stop it on an independent re-test that fails to reproduce the condition.

08

DIAGNOSTICS, NOT A PROJECT PLAN
Executive Diagnostic: Assessing Readiness for Interdiction

The fastest way to tell whether a programme has crossed from observation into interdiction is to ask a few questions and listen to the shape of the answer. The first set is for your own team, the second for anyone selling you a platform or a service. In both, the weak answer describes activity and the strong answer describes evidence.

Ask your own team

THE QUESTION	A WEAK ANSWER SOUNDS LIKE	A STRONG ANSWER SOUNDS LIKE
How long did our last five critical exposures stay exploitable?	“All closed inside SLA.”	“Median of X days, with the slowest at Y. The largest single interval was waiting for approval to act.”
Which of our open criticals have been proven exploitable here?	“They are all rated critical by the scanner.”	“Twelve of ninety carry reproduction evidence from our own environment. The rest are queued for validation, and are not being treated as critical until they are.”
What can we turn off, isolate or revoke within four hours, and who signs?	“It depends. We would raise an emergency change.”	“Here is the pre-authorised menu, with blast radius, approver and rollback defined for each option.”

Where are our model endpoints and agent tool servers, and what can they reach?	"The AI team owns that."	"They sit in the asset register with their authentication posture and reachable scope recorded, and they are re-tested on the same cadence as anything else internet-facing."
If our discovery tooling doubled its findings tomorrow, what would improve?	"Coverage, and our risk score."	"Nothing, because remediation capacity is the binding constraint. That is why the next increment of budget goes to validation and interdiction, not to discovery."

Ask a platform or a partner

THE QUESTION	A WEAK ANSWER SOUNDS LIKE	A STRONG ANSWER SOUNDS LIKE
How do you establish exploitability?	"We enrich with threat intelligence and scoring to prioritise."	"We attempt it safely, in your environment, hand you the evidence and the reproduction path, and publish our validated-true-positive rate."
What happens after a finding is confirmed?	"We integrate with your ticketing system."	"We propose a specific reversible control, with blast radius, approver and rollback, and we re-test the closure."
How do you handle AI and agent infrastructure?	"It is on the roadmap." Or silence.	"We discover model endpoints and agent tool servers, assess their authentication posture, and test what an injected instruction can reach."
What do you measure success by?	"Findings surfaced, coverage, risk score reduction."	"Reduction in the time an exploitable condition stayed open, evidenced per exposure."
Where does AI sit in your operating model?	"AI-powered" used as an adjective, unexplained.	"Here is which decisions it makes, which it recommends, what the guardrails are, where a human signs, and what rolls back when it is wrong."

IF YOU TAKE ONE THING FROM THIS SECTION

Both tables test the same distinction. An answer built on activity tells you what the team or the vendor has been busy with. An answer built on evidence tells you what is no longer exploitable, and how long that took. Only the second holds up in front of a board, an auditor or an attacker.

Operationalizing Preemptive Exposure Management with NST Assure

NST Assure is a continuous threat exposure management platform and service built around the discovery-to-neutralisation interval, with visibility as an input rather than a deliverable.

- | **Continuous discovery in environmental context.** Attack-surface mapping grounded in what is reachable from where an adversary stands, not a generic scan projected onto an asset list.
- | **Validated exploitability as the prioritisation input.** Ranking driven by evidence that an exposure is workable in your environment and by what it would reach next, rather than by severity alone.
- | **Support for the move from finding to action.** Interdiction options presented with the context a change owner needs to approve them quickly and reverse them safely.
- | **Domain depth where generic tooling is weakest.** Identity and non-human identity access paths, cloud configuration, supply chain, and emerging AI and agent infrastructure.
- | **A service-led operating model.** Platform capability paired with the analyst expertise to interpret results, defend a judgement call and drive a remediation outcome through an organisation with other priorities.

We treat preemptive exposure management less as a category to claim than as a stricter standard for running exposure management honestly.

WHY THIS IS A TRUST QUESTION, NOT ONLY A SECURITY ONE

Trust is not conferred once, at certification. It is tested continuously: by customers, by regulators, by third-party risk teams, and by attackers.

Every claim about your security posture should rest on current, testable evidence rather than a point-in-time attestation. Exposure management produces that evidence, but only if it validates and acts. An unvalidated finding proves nothing. A closed, evidenced window proves something specific: that when a real, exploitable condition appeared, you found it, proved it and removed it in a measurable time.

MTTN, reported with its tail, answers a question the board and the customer are both already asking.

Executive Takeaways

- | Assume exploitation precedes your awareness. Design for that case, not the exception.
- | Treat un-validated findings as a cost, not an asset, and move the next increment of automation to the closing half of the process.
- | Build the interdiction menu before you need it. Governance friction is the largest interval in most organisations, and the cheapest to fix.
- | Measure the window; report the tail.

CONTACT

To discuss a baseline MTTN assessment against your own environment, contact your NetSentries representative.

EVIDENCE BASE

Gartner: The Future of Exposure Management Will Be Preemptive, Driven by Autonomous Interdiction (Feb 2026); Market Guide for Adversarial Exposure Validation (Mar 2026); Top Funded Startups for Preemptive Exposure Management (Apr 2026); Security and Risk Management Summit guidance and AI-agent governance predictions (Apr to Jun 2026).

Threat and incident telemetry: CrowdStrike 2026 Threat Hunting Report (Aug 2026) and 2026 Global Threat Report (Feb 2026); IBM Cost of a Data Breach 2026 (Jul 2026); Verizon 2026 DBIR (May 2026); Mandiant M-Trends 2026 (Mar 2026) and Google Threat Intelligence Group (May 2026); Cisco Talos IR Trends Q2 2026 (Jul 2026); Sysdig (Jul 2026); named incidents from published reporting (Jul 2026).

Government, standards and survey: US NSA Model Context Protocol: Security Design Considerations (May 2026); NIST AI Agent Standards Initiative (Feb 2026); UK AI Security Institute cyber capability evaluations (Apr 2026); Anthropic Claude Opus 5 system card and containment disclosures (Jul 2026); SANS 2026 Identity Threats and Defences Survey (Apr 2026); Prophet Security State of AI in the SOC 2026 (Aug 2026); Cyentia Institute remediation-capacity research.

REPRESENTATIVE VENDOR



2026 Gartner® Market Guide for Adversarial Exposure Validation

NST ASSURE™ PEM · CTEM · AI-Driven AEV

GARTNER PEER INSIGHTS™



Highly Rated

NST ASSURE™
PREEMPTIVE EXPOSURE MANAGEMENT

✉ info@nstcyber.ai

🌐 www.nstcyber.ai

📍 99 S Almaden Blvd,
Suite 600, San Jose
CA 95113

📍 G011, Technohub,
DTEC, Dubai Silicon
Oasis, Dubai

📍 No 185/7, 2nd Floor, Chandra
Plaza, 8th F Main, 3rd Block
Jayanagar Bangalore, India

About NST Cyber

NST Cyber, a premier provider of advanced exposure assessment and adversarial validation solutions for large enterprises, utilizes AI and ML in NST Assure to proactively identify and prioritize exposed assets and vulnerabilities. This approach enables organizations to mitigate risks based on potential exploitability and threat patterns before malicious actors can act.