

PICUS

BLUE REPORT™

2026

THE STATE OF
THREAT EXPOSURE MANAGEMENT



Table of Contents

Introduction 3

Executive Summary 4

Key Findings 5

Key Recommendations 7

Picus Meets AI-Speed Attacks
with AI-Speed Validation 9

Methodology 10

Scoring Legend 11

Overall Prevention and Detection
Effectiveness Performance 12

Real-World Performance of
Cybersecurity Products 15

Uncovering Critical Defensive Gaps with
Autonomous Penetration Testing..... 17

Detection Rule Effectiveness 19

Performance by Industry 21

Performance by Region 25

Performance by Attack Vector 28

Performance by MITRE ATT&CK® Tactics 30

Performance by MITRE ATT&CK® Techniques 32

Performance by Threat Group 34

Spotlight on Ransomware Attacks 36

Spotlight on Vulnerabilities 38



Prevention Score
went from 62% to



69%



Log Score
went from 54% to



58%



Alert Score
went from 14% to



14%

Introduction

At a Glance

Now in its fourth year, the **Blue Report** continues to measure the real-world effectiveness of enterprise prevention and detection capabilities by analyzing how security controls perform against adversary behavior, not in a lab, but in production. Developed by Picus Labs, this annual study draws on **over 338 million attack simulations** executed across real-world customer environments in the first half of 2026, providing a comprehensive view of how security products and configurations perform, and what adversaries can actually achieve once they get inside.

We've designed The Blue Report 2026 to serve as **a practical guide for security teams and decision-makers aiming to mature their security posture through Continuous Threat Exposure Management (CTEM).**

By identifying blind spots, validating assumptions, and fine-tuning defenses based on the latest adversary behaviors, organizations can reduce uncertainty, better prioritize their resources, and build a more resilient cybersecurity foundation.

What's New for 2026

- This year's analysis connects tactical-level prevention, technique-level prevention, and live attack-path execution into one finding: defenses tuned for conspicuous activity are being outmaneuvered by adversaries who purposefully stay quiet.
- With adversaries now weaponizing new vulnerabilities in hours, the report reflects how Picus users leverage AI-driven validation to compress the cycle of testing, remediation, and re-validation **from weeks into minutes.**



Executive Summary



Post-Compromise
Prevention Rate



IOC-Based
Prevention Rate



Stealth
Prevention Rate

The Blue Report 2026 describes a year of **genuine recovery shadowed by a persistent structural weakness**. Average prevention effectiveness rose from 62% to 69%, reversing last year's decline and returning to its 2024 peak, while logging climbed to a four-year high of 58%.



Security controls degrade without validation and recover when they're put back to work.
What neglect lets slip, validation can restore.

But recovery at the boundary often masks a vulnerable interior. For the first time, the report measures what an adversary accomplishes after gaining authenticated access, using Autonomous Penetration Testing. In fact, **the Post-Compromise Prevention Rate was just 37%, far below the 69% overall score**. Prevention splits along a sharp line: traditionally "loud" actions like lateral movement and privilege escalation were caught, while quiet, stealthier actions tended to run almost unopposed. **The Stealth Prevention Rate, covering discovery and collection, reached an alarmingly low 10%**. Modern attackers can still map your domain, harvest your sessions, and read your credentials with little resistance.

That fault line runs throughout. Stealth was one of only two tactics to decline in prevention, and the least prevented technique was Impair Command History Logging (T1562.003) at 1%. **The same weakness shows in malware defense, where the IOC-Based Prevention Rate fell to 50%**, down a full 21 points over two years as indicator-based detection continues to lose ground. The reason is structural: VirusTotal alone receives nearly 2 million new files a day, a volume of fresh indicators no signature-based approach can possibly keep pace with. IOC-based prevention will always trail the threat, which is why it must be reinforced with TTP-based prevention that targets adversarial behaviors underneath the indicators, the behaviors that stay constant even as hashes and infrastructure churn. These results corroborate **The Red Report 2026** findings that adversaries are leaning hard into stealth.

Detection remains the most stubborn gap. This year, logging improved, but the alert score held flat at 14%, making the log-to-alert gap look more like a detection engineering problem and less like a collection one.

The Blue Report 2026 reinforces the fact that static, human-paced defenses are no match for adaptive threats operating at machine speed. In the post-Mythos era, adversaries are weaponizing new vulnerabilities in roughly eight hours. To stay ahead, organizations must adopt Adversarial Exposure Validation and AI capabilities into their CTEM programs to move beyond assumptions, prioritize with confidence, and respond at the same daunting speed as their attackers.



Key Findings

This year's report reveals a cybersecurity landscape in partial recovery, where broad gains in prevention and logging are tempered by persistent, and in some cases deepening, failures against the quiet techniques adversaries increasingly employ. This year's analysis shows both what continuous validation makes possible and where its absence leaves organizations exposed. Below are some of the report's most significant findings:

Defenses Are Strong at the Perimeter but Soft on the Inside

1

After falling from 69% to 62% last year, the average prevention score bounced back to 69% in 2026. However, Autonomous Penetration Testing tells a different story than the overall prevention score. Averaged across the individual attacker actions tested, only 37% were blocked. The interior is not uniformly weak, though; here too, prevention splits along a clear line. Actions that execute code or move between machines were blocked at high rates, while again, quieter actions ran almost unopposed. Organizations have hardened the noisy post-compromise steps, but an attacker inside these environments can still map the domain, locate high-value assets, and read credential material with very little resistance.

Threat Group and Ransomware Prevention Deteriorate Broadly

3

Prevention effectiveness declined against nine of the ten least prevented threat groups, with Infy collapsing from 56% to 28% to lead the list. Ransomware moved in the same direction: every one of the ten least prevented families scored 38% or lower, with Play falling from 50% to just 13% to become the hardest strain to prevent. In each case, these are well-documented adversaries with published tradecraft, yet prevention fell even as the overall score rose, confirming that generic control improvement does not translate into more effective coverage against specific, evolving actors.

Adversaries Are Leaning Into Stealth

2

The Red Report 2026 found adversaries increasingly favoring stealth-oriented techniques, and this year's Blue Report corroborates that shift from three separate directions. At the tactic level, Stealth was one of only two tactics to decline in prevention. At the technique level, the least prevented behaviors are led by stealth: Impair Command History Logging (T1562.003) at 1% and Signed Script Proxy Execution (T1216) at 9%. And in Autonomous Penetration Testing, the actions that succeeded most were precisely the low-noise ones, while most loud executions were caught. The same fault line appears in control testing, technique-level prevention, and live attack-path execution, corroborating that defenses tuned for conspicuous activity are being outmaneuvered by adversaries who keep their heads down.

Malware Prevention Falls Again as IOC-Based Detection Loses Ground

4

Preventing malware download scenarios declined to 50%, down from 60% last year and 71% in 2024, a 21-point slide over two years against one of the oldest and best-understood attack vectors. The cause is structural. With millions of new malware samples discovered each day, the churn of fresh indicators is more than any signature-based approach can keep pace with. As adversaries rotate infrastructure and repack payloads, IOC-based prevention will always trail the threat. It must be reinforced with TTP-based prevention that targets the adversary behaviors beneath the indicators, which persist even as hashes and infrastructure change.



Logging Improves, but Alerts Stay Flat

5

The log score rose from 54% to 58%, the highest level recorded in four years of Blue Report data. The alert score, however, held unchanged at 14%. Organizations are collecting more telemetry than ever, but they're not converting it into action any faster. Today, fewer than 1 in 7 simulated attacks produces an alert, and the widening gap between what's logged and what is alerted remains the single most persistent weakness in enterprise detection.

Endpoint Security Improves as the Assume-Breach Mindset Takes Hold

7

Prevention against endpoint scenarios rose to 83%, the third consecutive year of gains, and Privilege Escalation was the most improved tactic of the year, climbing 24 points up to 79%. Together, these show organizations increasingly building defenses around the expectation that attackers will get in, focusing on hardening the post-compromise layer where lateral movement and escalation occur. The assume-breach mindset is no longer aspirational language; it is increasingly present in the numbers.

Industry and Regional Leadership Proved Perishable

9

Last year's strongest sectors regressed while last year's weakest recovered. Transportation rose 29 points to 79%, while Education collapsed 30 points to 40%, now the least protected tracked industry. Regionally, South Asia climbed from last place to a share of first at 71%, while North America fell to the lowest prevention score of any region. Swings this large in a single year reflect how quickly both the threat landscape and organizations' own internal environments are changing, as new adversary techniques, infrastructure shifts, tooling updates, and configuration drift steadily erode a posture that was strong months earlier. Strong performance is rented, not owned, and it lasts only as long as the validation behind it continues.

Detection Rule Failures Shift to Performance from Log Issues

6

For the first time, performance issues are the leading category of detection rule failure, at 49% of all issues, more than doubling from 24% a year ago, while log collection issues fell to 41% from 50%. This shift is expected: as organizations expand telemetry coverage, the bottleneck moves downstream from collecting data to processing it efficiently. But the remaining log collection issues are the more dangerous half of the picture, because they represent silent failures where the behavior is never captured and no alert fires. The pattern points to the same conclusion throughout this year's report: your priority now should be detection engineering.

The Least Prevented Vulnerabilities Cluster Around the Same Weakness Types

8

Every one of the ten least prevented vulnerabilities was blocked in less than 25% of simulations, and the list is not a random scatter of exotic flaws. It concentrates on recurring weakness classes across ubiquitous software: memory-safety and improper-input-handling flaws in core libraries and utilities (OpenSSL, GNU InetUtils, 7-Zip, WinRAR), and local privilege escalation in operating system components (Linux libblockdev, Windows WinSock, macOS Sequoia). The clustering by weakness type, rather than by product, indicates that the gaps are systemic to how similar vulnerability classes are handled.



Key Recommendations

The findings of **The Blue Report 2026** highlight the growing need for organizations to keep validation, detection analytics, and control hardening at the center of their security strategy.

This year's recovery in prevention proves that gaps close when they are tested and tuned, whereas the areas that regressed show what happens when that discipline lapses. Based on our analysis, we recommend taking the following actions to reduce exposure and improve your threat readiness:



Validate Exposure, Not Just Inventory

Shift from simply knowing where exposures exist to **proving which ones matter**. Adopt the adversarial exposure validation approach that goes beyond asset and vulnerability discovery to confirm exploitability. This enables faster, more informed prioritization and reduces the noise and wasted effort created by purely theoretical risks.



Harden the Interior Against Quiet Actions, Not Just Loud Ones

Validate defenses against domain enumeration, session and file share discovery, and credential material read from memory and the registry as rigorously as against lateral movement and privilege escalation. As with malware delivery, the answer is behavioral detection that triggers on what an action does rather than on which tool signature it matches, since the same credential-theft tool slips past controls simply by using a less famous method.



Simulate Complete Ransomware and APT Kill Chains

With every leading ransomware family scoring 38% or lower and prevention failing against nearly every top threat group, your coverage can't afford to rely on last year's playbook. **Simulate full, up-to-date kill chains for the ransomware families and threat actors most relevant to your sector and geography,** including initial access, defense evasion, data exfiltration, and encryption or extortion, and confirm that your controls will interrupt them at multiple points.



Move Malware Defense Beyond Indicators to Behavior

Two consecutive years of decline in malware download prevention, now down to 50%, show that indicator-based detection is losing effectiveness as adversaries rotate infrastructure and repack payloads. **Shift weight toward behavior-based detection and sandboxing that evaluate what a payload does rather than whether it matches a known indicator,** and validate these controls regularly against current delivery techniques.



Close the Log-to-Alert Gap Through Detection Engineering

A 58% log score set against a 14% alert score means most new telemetry never translates to action. **Treat detection content as something built, measured, and maintained:** write rules against current adversary behavior, test that they fire, tune them to cut noise, and re-validate them as infrastructure and threats change. The same continuous-testing discipline that restored prevention is the most direct path to an alert score that finally moves in the right direction.



Prioritize Detection Analytics and Investigate Silent Logging Failures

This year's detection rule improvements rightly focused on performance, but **the remaining log collection issues deserve equal attention because they fail silently, capturing nothing and triggering no alert.** Regularly validate log source health, coverage, and coalescing behavior, and pair that with stronger detection analytics so that captured telemetry is actually correlated into high-fidelity alerts. **Remember, collection without analytics is visibility wasted.**



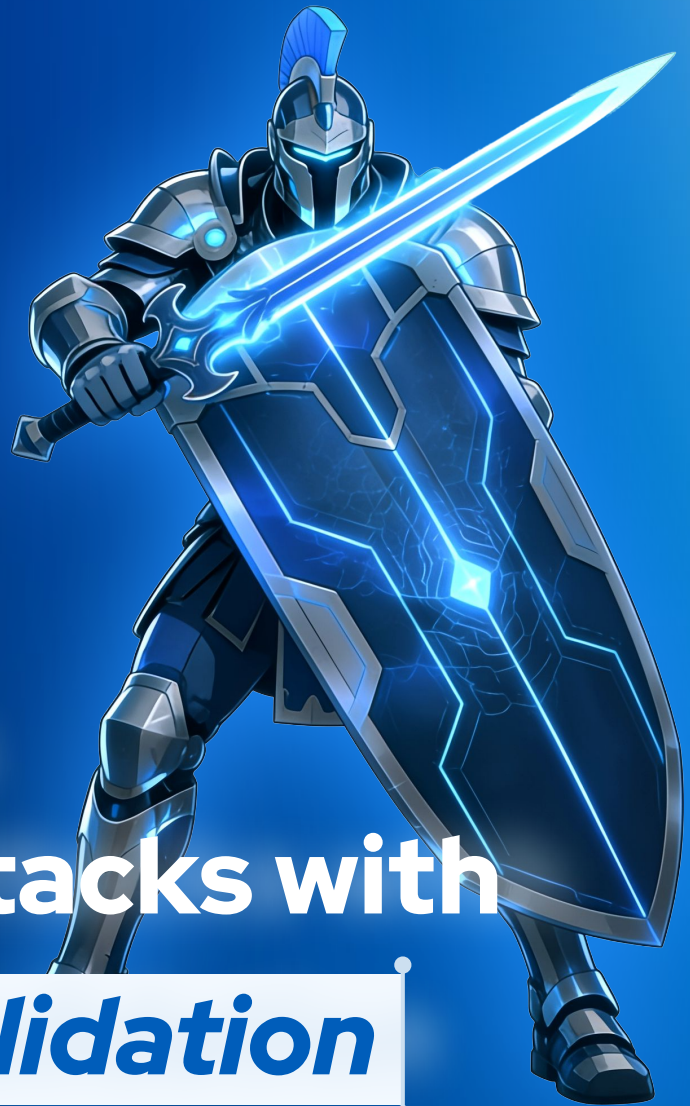
Defend the Post-Compromise Layer and Maintain the Assume-Breach Gains

The improvement in endpoint and privilege escalation defense shows the assume-breach mindset is actually paying off. Continue validating lateral movement, privilege escalation, and other post-compromise behaviors, and **never assume your endpoint hardening is finished, since macOS prevention regressed from 76% to 61% precisely because last year's gains were not maintained.**



Prioritize Stealth and Evasion Techniques in Validation

With this year's Red and Blue Reports both pointing to a decisive adversary shift toward stealth, defenses must follow. **Validate controls specifically against evasion behaviors such as command history impairment, signed binary proxy execution, encoded command and control, and protocol downgrade,** and strengthen the behavioral analytics needed to catch the deliberately low-noise activity that static, signature-led controls keep missing.



Picus Meets AI-Speed Attacks with *AI-Speed Validation*

Picus Autonomous Exposure Validation Platform brings together Breach and Attack Simulation, Autonomous Penetration Testing, and Exposure Validation, with Numi AI and Picus Swarm as the intelligence and orchestration layer. Together they help organizations of all sizes prove what attackers can actually exploit, prove what their defenses stop, and turn every exposure into a defensible decision, continuously and at machine speed.

In the post-Mythos era, adversaries weaponize a new vulnerability in roughly eight hours, and no human-paced program can keep up. Picus customers close that gap with the same machine speed. Using AI Threat Builder, they turn raw threat intelligence into a runnable, ATT&CK-mapped attack simulation in about nine minutes, down from days, red teaming at a tempo manual teams cannot match. Paired with one-click re-test, the full cycle of validation, remediation, and re-validation collapses from weeks into minutes.

With Picus, security leaders move beyond severity-led, scheduled validation to autonomous, evidence-backed readiness, matching AI-speed attackers with AI-speed validation while a named human stays on every risk-acceptance decision. Instead of drowning in findings they cannot act on fast enough, they consistently prove which attacks would succeed and close those doors first.



Methodology

The findings in this report are based on the results of simulated attack scenarios executed by Picus Security customers from January to June 2026. The data has been anonymized and aggregated from over 338 million attack simulations. Research and analysis were completed by Picus Labs and Picus Data Science teams.

Definitions

Prevention Effectiveness evaluates an organization's ability to block potential cyberattacks through its security controls. This metric is the percentage of successfully prevented attacks out of all simulated attacks executed. For example, an effectiveness score of 80% means that 80 out of every 100 simulated attacks were effectively prevented.

A high prevention effectiveness score indicates strong security controls that significantly lower the risk of successful breaches. Conversely, a low score highlights gaps in an organization's security measures, suggesting the need for security teams to conduct a thorough review and enhance their controls.

Detection Effectiveness assesses an organization's capability to identify potential cyber threats through its existing security controls. This report uses two key indicators for evaluating detection performance: Log Score and Alert Score.

- **Log Score:** This measures the percentage of simulated attacks where the attackers' behavior was logged. A higher log score demonstrates the efficacy of monitoring controls like SIEMs in capturing a large volume of events and identifying threat indicators. Effective logging is crucial for maintaining a comprehensive security posture and understanding attack patterns.
- **Alert Score:** This indicates the percentage of simulated attacks that generate alerts. High alert scores are crucial for ensuring that security teams are promptly informed of any and all threats, enabling them to take immediate action to neutralize potential risks. Alerts serve as critical triggers for initiating a timely and effective response to attacks.



Scoring Legend

Results are color-coded and categorized into five distinct levels of threat exposure management: Inadequate, Basic, Moderate, Managed, and Optimized (see table below).

This classification provides a clear, visual representation of an organization's cybersecurity effectiveness, facilitating easy benchmarking and identification of areas for improvement.

90-100%
Optimized

Organizations with optimized security controls continuously monitor, refine, and update them to keep up with the evolving threat landscape and maintain their edge in exposure management.

70-89%
Managed

Managed security controls offer a high level of protection against a wide range of threats, significantly reducing the risk of successful attacks. Organizations at this level should maintain their strong security posture, regularly assess the effectiveness of their controls, and address identified gaps in exposure management.

40-69%
Moderate

Moderate security controls provide a reasonable level of protection against various threats. Organizations at this level should continue to refine their security controls and consider additional measures to further reduce their threat exposure.

20-39%
Basic

Basic security controls offer limited protection against a narrow range of threats. Organizations at this level should invest in enhancing and expanding their security controls to achieve a more effective threat exposure management program.

0-19%
Inadequate

Inadequate security controls provide almost no protection to minimal protection against threats, leaving the organization highly vulnerable to attack. At this level, only a few basic security measures are in place. Organizations with this level of exposure need to urgently review and improve their security posture.



Overall Prevention and Detection Effectiveness Performance

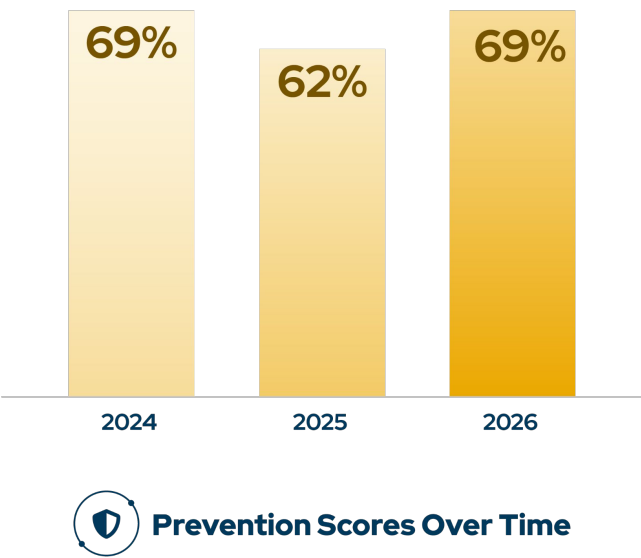
This year's report reflects a landscape in recovery, but an uneven one. Prevention effectiveness rebounded strongly, logging coverage reached its highest recorded level, and yet alerting performance did not move at all. Taken together, the results show that many organizations acted on last year's warning signs, while also confirming that the hardest part of the detection pipeline, turning telemetry into timely alerts, remains unsolved. Closing that gap depends on detection engineering: the ongoing work of writing, testing, and tuning the rules that turn logs into alerts. It is not a one-time build but a continuous effort, and the organizations that treat it that way are the ones that convert visibility into detection.

Prevention Effectiveness

In 2026, the prevention effectiveness score rose to 69%, up from 62% in 2025. This recovery erases last year's seven-point decline and returns prevention performance to the peak previously recorded in 2024, continuing the longer arc of improvement from 59% in 2023. The rebound suggests that organizations that slipped behind on control maintenance last year have re-engaged with validation and tuning, closing gaps that had opened through configuration drift and unaddressed adversary evolution.

The four-year trend shows that prevention effectiveness is not a plateau that organizations reach and hold. It is a level that must be actively maintained, and the difference between a rising year and a falling year is the discipline of continuous exposure validation and real-world attack simulation.

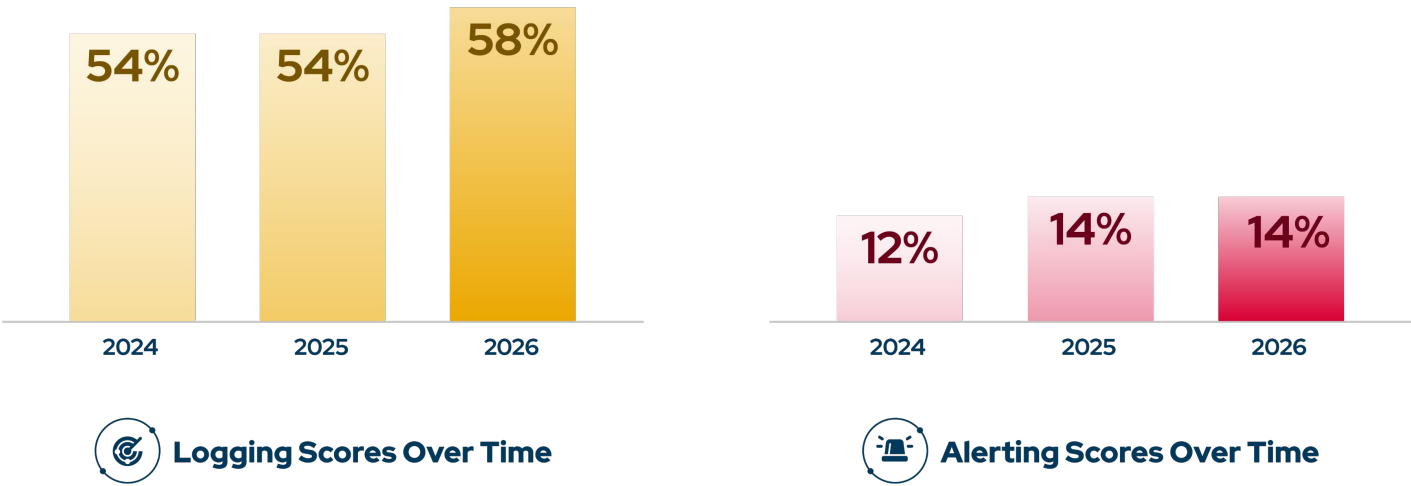
At 69%, organizations are blocking roughly two out of every three simulated attacks. That is meaningful protection, but it also means nearly a third of attacks still succeed against production defenses. The goal for the year ahead is not to celebrate the recovery but to consolidate it, so that 2026's gains do not become 2027's regression.





Detection Effectiveness

In 2026, detection effectiveness presents a split picture. **The log score rose from 54% to 58%**, the first improvement in two years and the highest log score recorded across four editions of this report. After two consecutive years at 54%, the gain indicates real investment in telemetry pipelines, log source coverage, and visibility across enterprise environments.



The alert score, however, held flat at 14%. Fewer than 1 in 7 simulated attacks generates a meaningful alert. For most organizations, this level of performance is far below what is needed for timely detection and effective response, and the fact that it did not improve while logging did is itself the finding: more telemetry, on its own, does not produce more detection.

The widening gap between what is logged (58%) and what is alerted (14%) highlights systemic issues in detection engineering. Many organizations are still operating with detection pipelines that are prone to silent failure, either because telemetry never reaches correlation logic in usable form, or because detection rules fail to trigger on relevant behaviors. Collecting telemetry is not the same as understanding it, and logging alone does not guarantee that threats will be detected, prioritized, or acted upon.

To address these issues, organizations need to move beyond basic telemetry collection and invest in validating the entire detection lifecycle. This includes not only verifying log availability and quality but also testing whether detection rules are functioning as intended and whether they generate high-fidelity alerts in response to real-world attacker behaviors. Without these validations, organizations risk overestimating their detection capabilities and underestimating their exposure.

Addressing the Gaps

The value of a multi-year dataset is that it turns assertions about security into something closer to evidence, and this year's numbers settle a question the previous edition could only argue. Last year we claimed that controls degrade without validation. In 2026, prevention did not merely stop falling; it climbed seven points back to its previous peak. That reversal is difficult to explain by anything other than organizations re-testing controls that had drifted and fixing what the tests exposed. The degradation thesis now has a matching recovery thesis: what neglect lets slip, validation can restore.



Detection has to keep up with log collection.



The alert score is where that logic has not yet been applied. Logging rose four points to a four-year high, yet the share of attacks that produce an alert did not move. The two metrics diverged because they are governed by different work. Broader log coverage comes from pipeline and infrastructure investment, while better alerting comes from detection engineering, and a flat alert score is the signature of that second effort lagging the first. Telemetry sitting in a SIEM does not correlate itself.

That reframes what closing the gap means for the year ahead. It is less a tooling problem than a practice problem. Detection content has to be treated as something built, measured, and maintained, with rules written against current adversary behavior, tested to confirm they fire, tuned to cut noise, and re-validated as infrastructure and threats change. Organizations restored prevention by putting controls under continuous test. Extending that same discipline to detection rules is the most direct path to much improved alert scores.

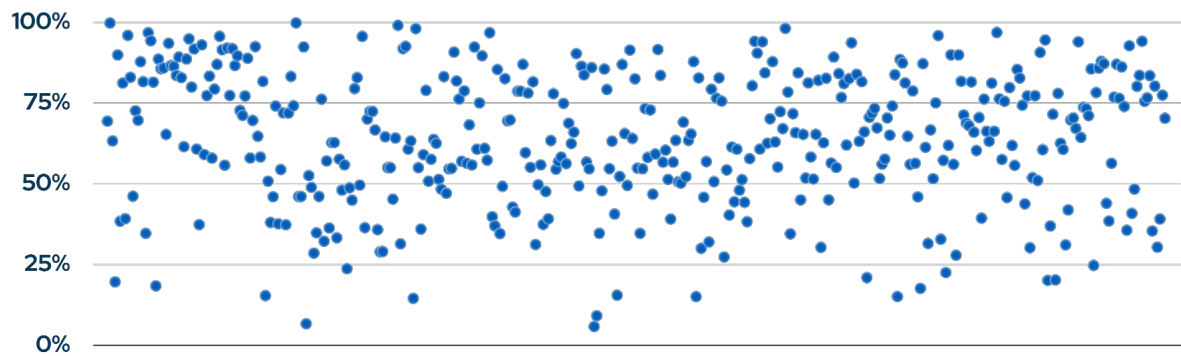


Real-World Performance of Cybersecurity Products

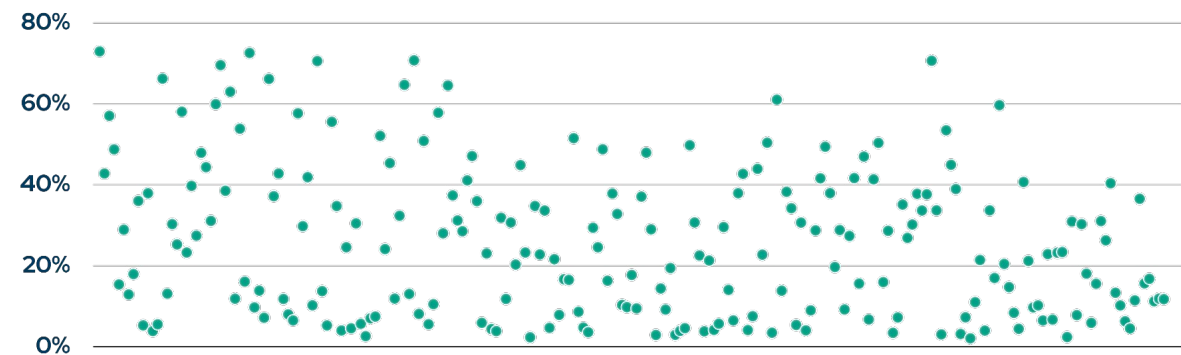
Cybersecurity products are often evaluated in controlled test environments using curated attack scenarios, tightly defined metrics, and ideal configurations. **Benchmarks like ATT&CK® Evaluations offer valuable insight into a product's potential, but they do not reflect the full complexity of production environments**, where configuration drift, integration friction, and operational constraints can significantly impact outcomes. This year's data tells the same story: the score a product earns in a lab is a ceiling, not a promise of how it performs in production.

A flawless lab score is a statement about the product, not about the environment it lands in. **Once deployed, even top-scoring tools protect inconsistently, less because of what they cannot do than because of how they are configured and kept up.**

Prevention Score



Detection Score





Our analysis shows that control degradation is more common than most teams realize. Security tools that were properly installed and configured at deployment may fail over time due to:

- **Configuration drift:** Policy changes, software updates, and integration rollouts gradually alter how tools behave.
- **Broken integrations:** A misconfigured log source or disabled alerting rule may render a detection pipeline ineffective without anyone noticing.
- **Operational complexity:** Each organization has a unique mix of systems, staffing levels, compliance requirements, and network architecture. These variables can significantly impact product effectiveness.
- **Dynamic threats:** The threat landscape evolves constantly. Products must be regularly tested and updated to keep up with new TTPs. Static signatures and stale configurations quickly become obsolete.

Teams often read a product's presence in the stack as confirmation that it functions. Deployment and effectiveness are different things. Absent validation, whether a detection rule fires, a prevention control blocks a live threat, or an alert ever gets escalated all remain open questions.

This gap between assumed and real performance opens exposure that no amount of budget alone can close, and it persists inside even the best-resourced programs. **Security tools return their intended value only when they are maintained continuously**, not treated as finished the day they go live.

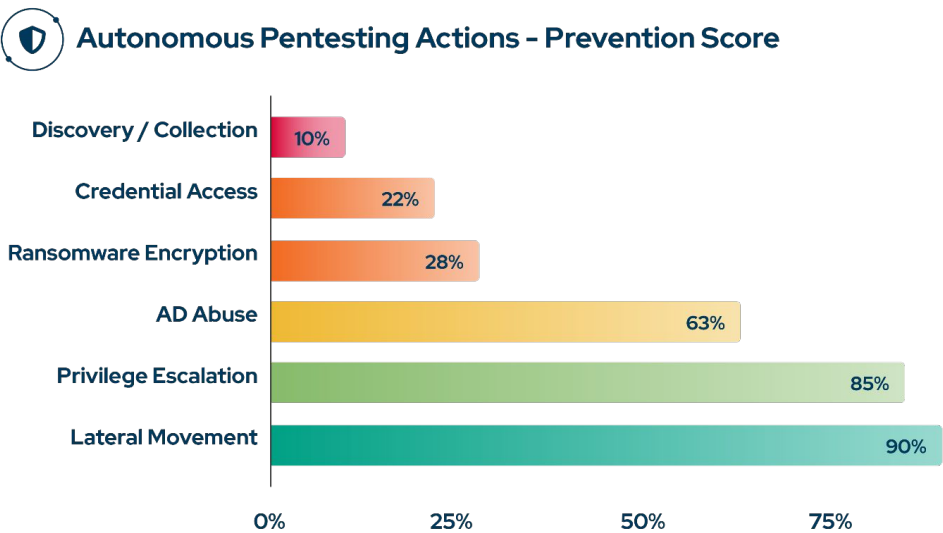
We recommend:

1. **Treat evaluations as guidance, not guarantees.**
MITRE ATT&CK® Evaluations and other independent tests offer valuable starting points. But they must be validated in your own environment, with your configurations, controls, and constraints.
2. **Expect degradation and plan for it.**
Even the best tools lose effectiveness without maintenance. Routine validation is the only way to detect control drift and silent failures before attackers do. This year's overall recovery shows the same mechanism working in reverse: tested controls improve.
3. **Make continuous validation part of your workflow.**
Use simulated attack scenarios to test both prevention and detection capabilities in context. Confirm that controls behave as expected, that alerts are generated, and that response processes activate when they should.



Uncovering Critical Defensive Gaps with Autonomous Penetration Testing

Autonomous penetration testing remains a vital capability for uncovering real-world weaknesses across enterprise environments. Leveraging the **Picus Autonomous Penetration Testing** product, organizations can emulate full-chain adversarial behavior, including credential compromise, lateral movement, and privilege escalation, without requiring red team resources or intrusive manual efforts. Where the overall prevention score measures how well controls stop attacks at the boundary, this module measures something harder: what an adversary can accomplish once already operating inside the environment as an authenticated user.



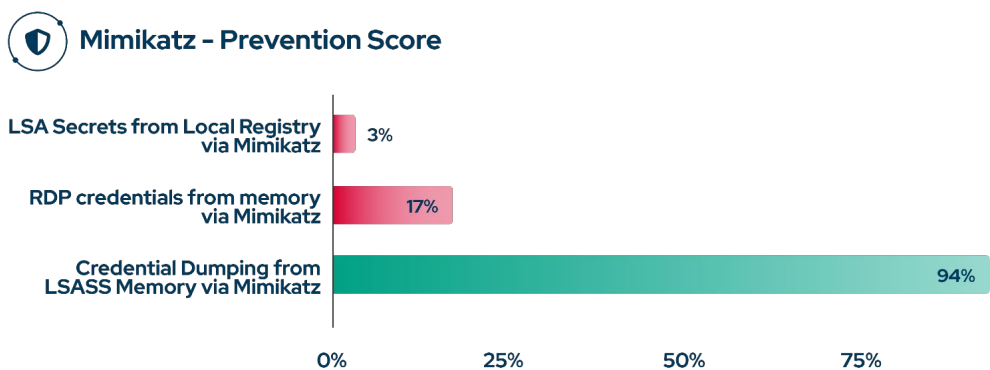
The results describe a soft interior. Averaged across the individual attacker actions tested, **only 37% were blocked**, meaning that inside the domain, **defenses stopped roughly one action in three** and let the other two proceed. That figure sits well below the 69% overall prevention score, and the gap between the two is the finding: perimeter and endpoint controls are recovering, but the assumption that an attacker who gets past them will still be contained does not hold up under test.

The most important pattern is not the average but the shape of the distribution behind it, because prevention splits almost cleanly along a single line: **loud actions are caught, quiet ones are not**. Actions that execute code or move between machines were blocked at high rates. Lateral movement through service execution techniques such as **Sharp-ServiceExec and SMBExec was prevented around 90%** of the time, privilege escalation through **UAC bypass techniques around 85%**, and credential reuse and Active Directory abuse, including Pass-the-Ticket and malicious certificate requests, around 63%. This is EDR doing its job: endpoint scenario prevention reached 83%, and Privilege Escalation was the most improved tactic of the year. The assume-breach investment is visible here in the actions defenders now reliably stop.



The actions that succeeded are the quiet ones. **Discovery and collection behaviors were the least prevented category by a wide margin**, blocked in only about 10% of attempts on a per-action basis, and several were not blocked at all. Domain enumeration through SharpHound, domain file share enumeration, session enumeration, and local file collection ran almost entirely unopposed. Credential material read passively from memory and the registry fared little better as a group, at roughly 22% prevention, with local registry secret extraction and bulk local admin credential validation checks blocked in well under 1% of attempts. An attacker inside these environments can map the domain, locate high-value assets, harvest sessions, and read credential material with very little resistance, and can do so before ever taking the noisier action that would trigger a control.

The Red Report 2026 documented adversaries deliberately shifting toward stealth. The Autonomous Penetration Testing data shows exactly that preference paying off inside real environments: the behaviors defenders miss are precisely the low-noise ones an evasion-focused adversary depends on.



One illustration captures the mechanism. **The same credential-theft tool produced sharply different outcomes depending on how conspicuous its method was.** The classic, heavily signatored path of dumping credentials directly from LSASS process memory was blocked in the large majority of attempts, while quieter variants of the same tool that read credentials from the registry or from other memory locations were blocked rarely or not at all. Controls are tuned to the well-known indicator, not to the underlying behavior, and adversaries who simply choose the less famous variant walk past them. This is the credential-access counterpart to the malware download finding: indicator-based detection is losing ground to behavior that stays just off the signatored path.

The practical implication is that infrastructure hardening cannot stop at execution and escalation.

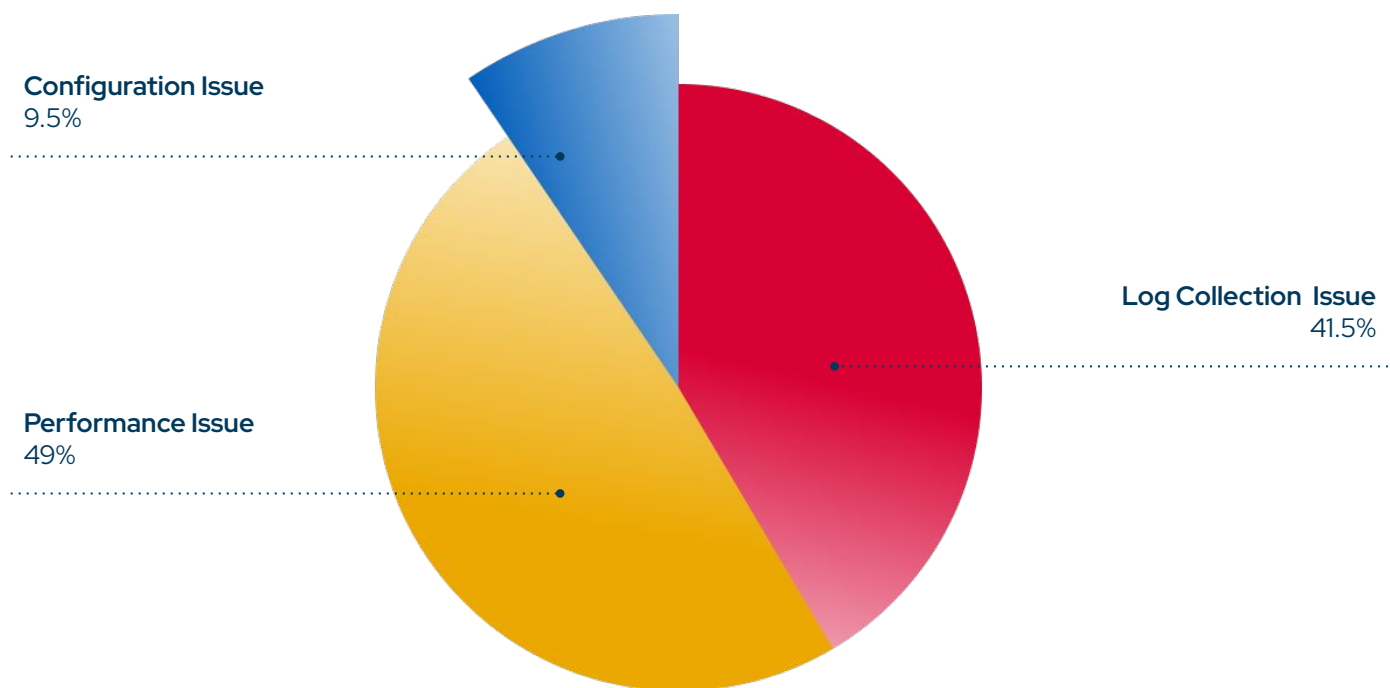
Organizations have made real progress on the actions that make noise, and that progress is worth protecting. But an interior in which reconnaissance and credential reading run unopposed gives an adversary everything needed to plan a precise, quiet path to the assets that matter, and to do it below the threshold of the controls that are working. Closing this gap means validating controls against discovery, collection, and passive credential access as rigorously as against lateral movement, and building behavioral detection that triggers on what an action does rather than on which tool signature it matches.



Detection Rule Effectiveness

Detection rules are the logic layer that decides whether telemetry becomes an alert. Written and maintained well, they let SIEM and EDR platforms surface meaningful signal from enormous volumes of events. Our analysis shows that this layer is where much of the alerting gap originates: this year's flat 14% alert score, set against a rising log score, is the outcome that broken and inefficient rules produce.

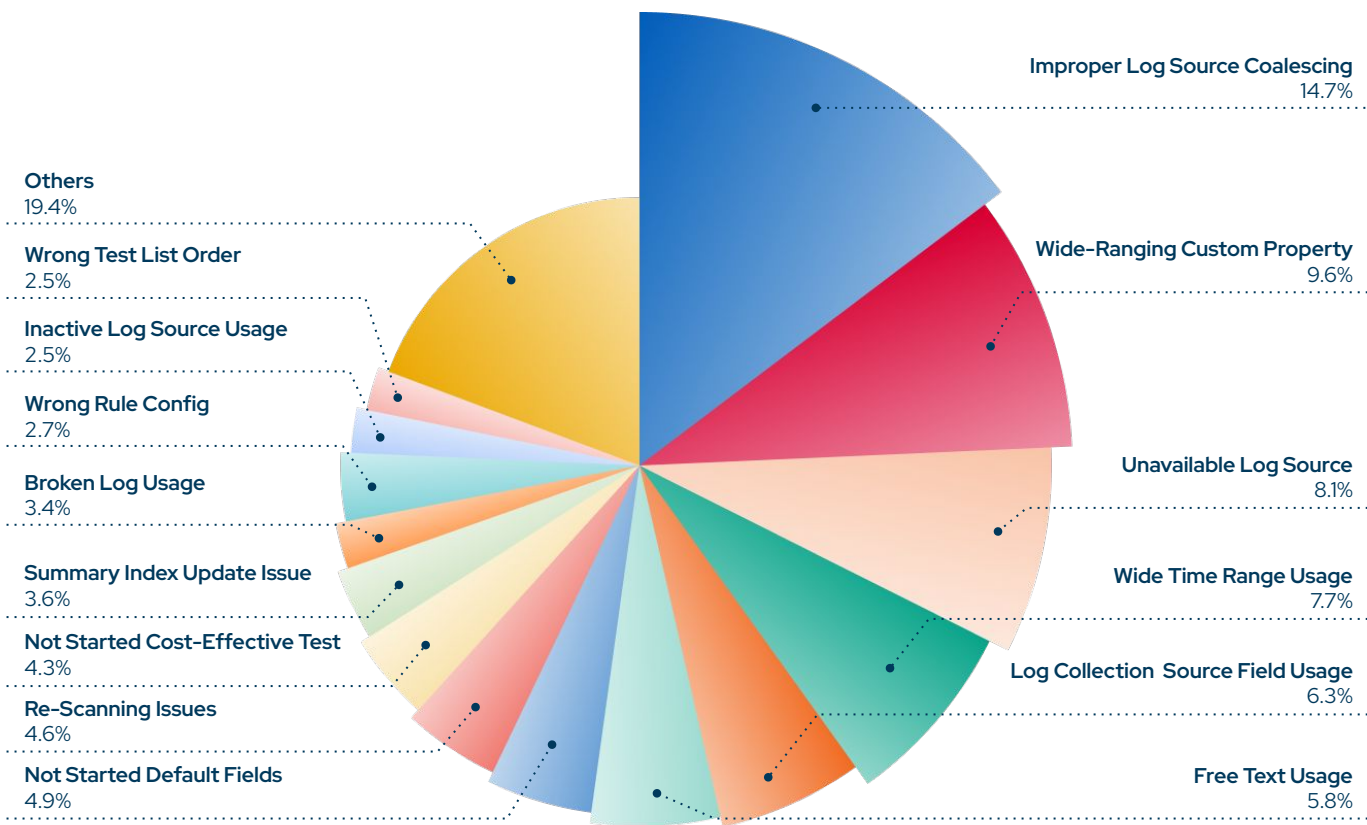
Common Issue Types in Detection Rules



This year's data also marks a shift in where rules fail. For the first time, **performance issues are the leading category of detection rule failure, at 49%** of all issues, up from 24% a year ago. Log collection issues, previously the dominant failure mode, fell to 41.5% from 50%, and configuration issues accounted for the remaining 9.5%, down from 13%. The movement is telling. As organizations expanded telemetry coverage, reflected in this year's higher log score, the bottleneck migrated downstream, from getting data in to processing it efficiently enough to alert on it. More logs did not fix detection. They relocated the problem.



Common Issues Affecting Detection Rule Effectiveness



The individual findings explain how performance became the dominant drag. **Wide-Ranging Custom Property Definition** rose to 10% of all issues, now the second most common insight overall, and Wide Time Range Usage (8%), Free Text Usage (6%), and Re-Scanning Issues (5%) round out a cluster of overly broad rule constructions. These configurations consume excessive processing power, slow query response, and produce imprecise matches that drive alert fatigue. A rule that is too expensive to run, or too noisy to trust, contributes nothing to detection even when it technically fires.

Log source problems, though no longer the largest category, remain the most concentrated individual failure. **Improper Log Source Coalescing was again the single most common insight, at 15%** of all issues, down from 21% but still the top line item. This occurs when event coalescing is enabled for log sources such as DNS systems, proxy servers, Windows servers, and endpoints, compressing or dropping events so that detection logic becomes partial, delayed, or entirely ineffective. Unavailable Log Source (8%), Broken Log Source (3%), and Inactive Log Source Usage (3%) reflect telemetry pipeline disruptions, misconfigured forwarding agents, and segmentation gaps. Each of these fell from last year, consistent with the improved log score, but together they still represent a visibility gap that detection teams must keep closing.

Configuration problems, the smallest category, are the quietest. **No Log Source Field Usage held steady at 6%** and Wrong Rule Config at 3%. These failures typically stem from drift between detection content and infrastructure changes, leaving rules that exist on paper but silently fail to trigger in production. Their low volume makes them easy to overlook, and their silence makes them dangerous.

The dataset for these statistics is sourced from Picus Breach and Attack Simulation (BAS), which uses a continuously updated checklist to identify and categorize over 50 recurring issues in detection rules. This year's breakdown reframes the detection challenge. Getting telemetry into the SIEM is no longer the primary obstacle; making rules run efficiently and fire precisely now is. Closing the log-to-alert gap therefore depends less on collecting more and more on continuous validation, fine-tuning, and pruning of the rules already in place, so that detection content stays accurate, affordable to run, and responsive to the behaviors that matter.



Performance by Industry

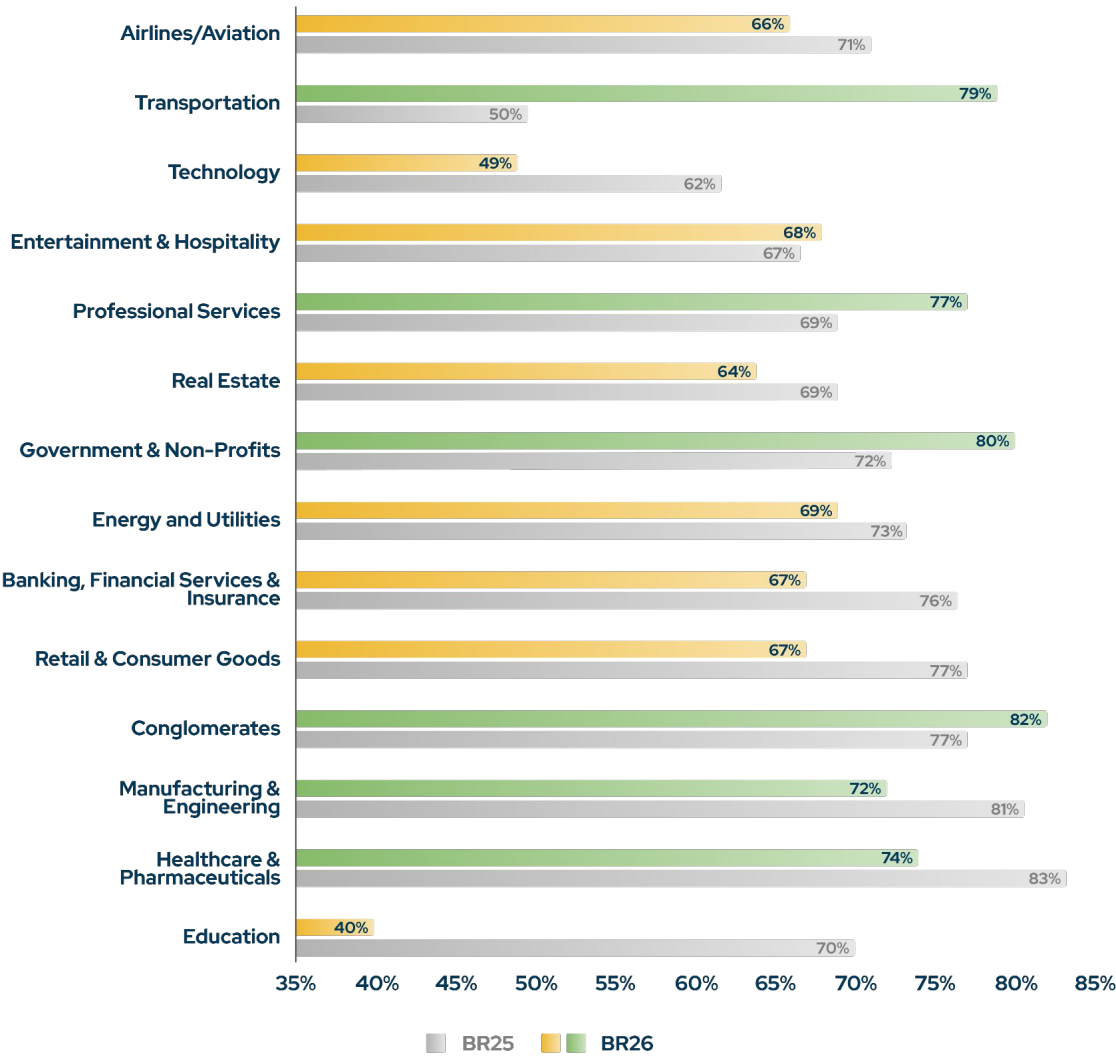
In this section, we examine the prevention and detection effectiveness across various industries, highlight significant changes year on year, and identify the most and least successful industries. Our analysis offers a comprehensive view of how different sectors are performing in their efforts to prevent and detect cyber attacks.

Prevention Effectiveness

In 2026, prevention effectiveness scores diverged more widely across industries than in any previous edition of this report. **Conglomerates led all sectors with an 82% prevention effectiveness score**, up from 77% last year, reflecting the benefits of mature, centrally governed security programs operating at scale. Government & Non-Profits followed closely at 80%, an eight-point improvement that builds on last year's steady 72% and suggests sustained investment in control validation across public sector environments.



Prevention Effectiveness Score by Industry





The most dramatic improvement belongs to Transportation. Last year's weakest sector at 50%, it rose 29 points to 79%, the largest single-year gain in the dataset. The turnaround demonstrates what focused remediation can achieve when validation findings are operationalized: gaps that left the sector exposed a year ago have been systematically closed. Professional Services also improved strongly, rising from 69% to 77%.

Last year's leaders, by contrast, gave ground. **Healthcare and Pharmaceuticals, the top performer in 2025 at 83%, fell nine points to 74%.** Manufacturing and Engineering dropped from 81% to 72%. Retail and Consumer Goods declined from 77% to 67%, and Banking, Financial Services, and Insurance (BFSI) fell from 76% to 67%. These are still moderate-to-managed scores, but the pattern is a warning: the sectors that led last year and eased off are precisely the ones regressing now. Strong performance is rented, not owned.

The middle of the table saw smaller movements. **Energy and Utilities slipped from 73% to 69%,** Entertainment & Hospitality edged up from 67% to 68%, Airlines/Aviation declined from 71% to 66%, and Real Estate fell from 69% to 64%.

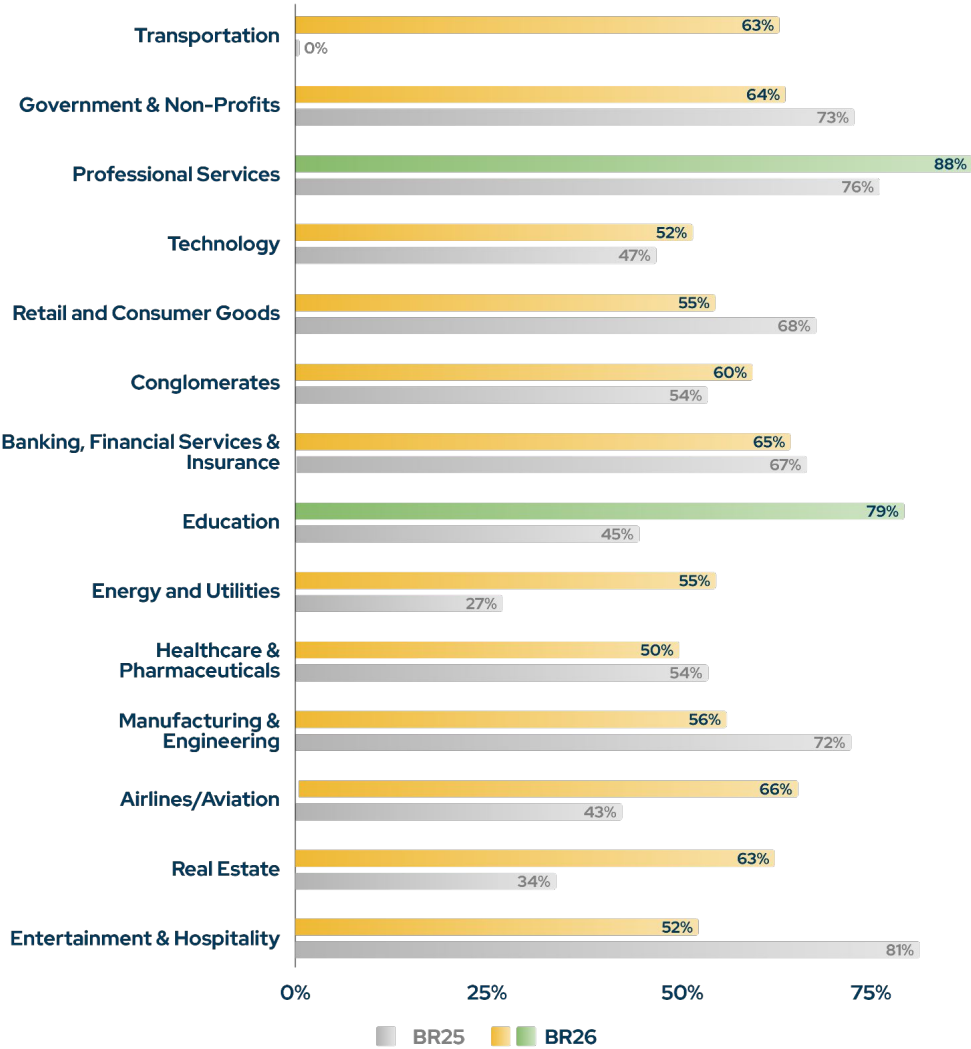
At the bottom, two sectors stand apart. **Technology fell from 62% to 49%,** extending last year's surprising underperformance and suggesting that innovation-led priorities continue to crowd out security control maintenance in a sector that should know better. Most alarming of all, Education collapsed from 70% to 40%, a 30-point decline that leaves it as the least protected industry by a wide margin. At 40%, Education environments fail to prevent six out of every ten simulated attacks. Resource constraints, sprawling and open network estates, and limited security staffing are longstanding challenges in the sector, but a drop of this magnitude signals systemic control failure that demands urgent review.



Detection Effectiveness

This year's findings reveal a continued disparity between logging and alerting performance across industries, highlighting the persistent challenge of transforming raw telemetry into actionable detection. Several sectors made major gains in log collection coverage, yet most are still struggling to generate timely and meaningful alerts.

 Log Score by Industry





Professional Services led all industries in log score with an impressive 88%, up from 76% last year. Its alert score also improved from last year's 5% to 13%, but the 75-point spread between logging and alerting remains the widest in the dataset, indicating that world-class visibility is still being squandered by underdeveloped correlation and alert logic.

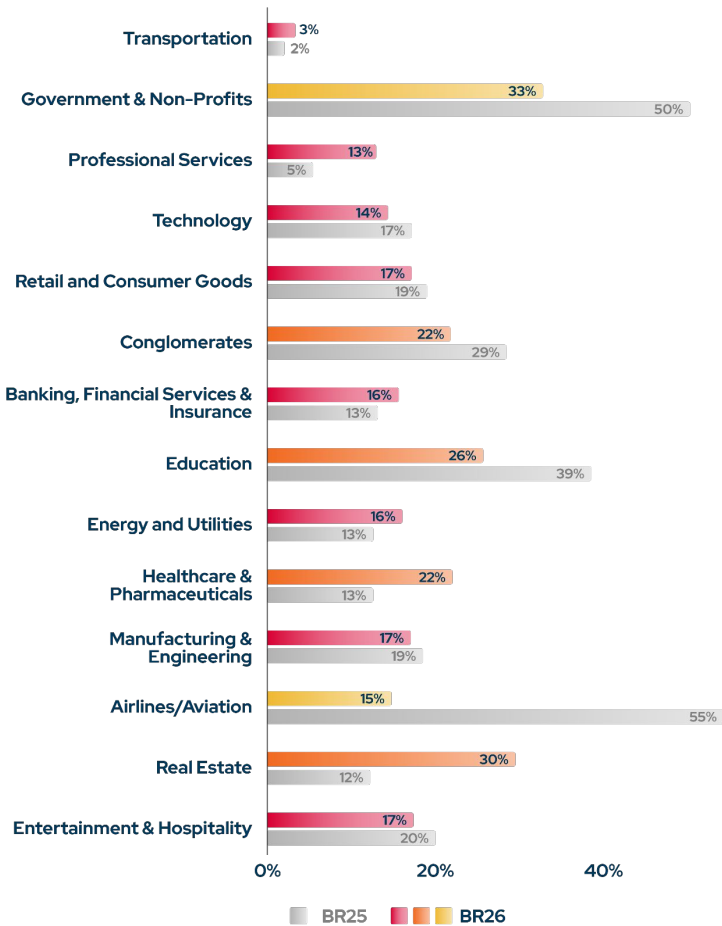
Education presents this year's most striking contradiction. Its log score surged from 45% to 79%, the second highest of any industry, even as its prevention score collapsed. Its alert score, however, fell from 39% to 26%. The sector can increasingly see attacks it cannot stop, and it alerts on fewer of them than it did a year ago.

Government & Non-Profits remained the alerting leader at 33%, though this represents a decline from last year's 50%, and its log score also slipped from 73% to 64%. The sector still demonstrates the most balanced detection performance overall, but the year-over-year softening suggests detection content is not being maintained at the pace it was built.

Real Estate quietly posted one of the strongest detection improvements of the year, with its log score rising from 34% to 63% and its alert score climbing from 12% to 30%, now the second highest of any industry. Airlines/Aviation improved its log score from 43% to 66%, but its alert score fell sharply from last year's outlier 55% to 15%, unwinding the unusual inversion noted in the previous edition.



Alert Score by Industry



Energy and Utilities doubled its log score from 27% to 55%, though alerting remains modest at 16%. Healthcare and Pharmaceuticals recorded the lowest log score of any industry at 50%, down from 54%, though its alert score improved from 13% to 22%. Entertainment & Hospitality, last year's log score leader at 81%, fell 29 points to 52%, the sharpest visibility decline in the dataset, with its alert score easing from 20% to 17%.

Among the remaining sectors, **BFSI held a managed-range log score of 65% with alerts at 16%**, Conglomerates improved to a 60% log score with a 22% alert score, Manufacturing and Engineering declined from 72% to 56% on logging with alerts at 17%, Retail and Consumer Goods fell from 68% to 55% with alerts at 17%, and Technology improved modestly to a 52% log score with a 14% alert score.

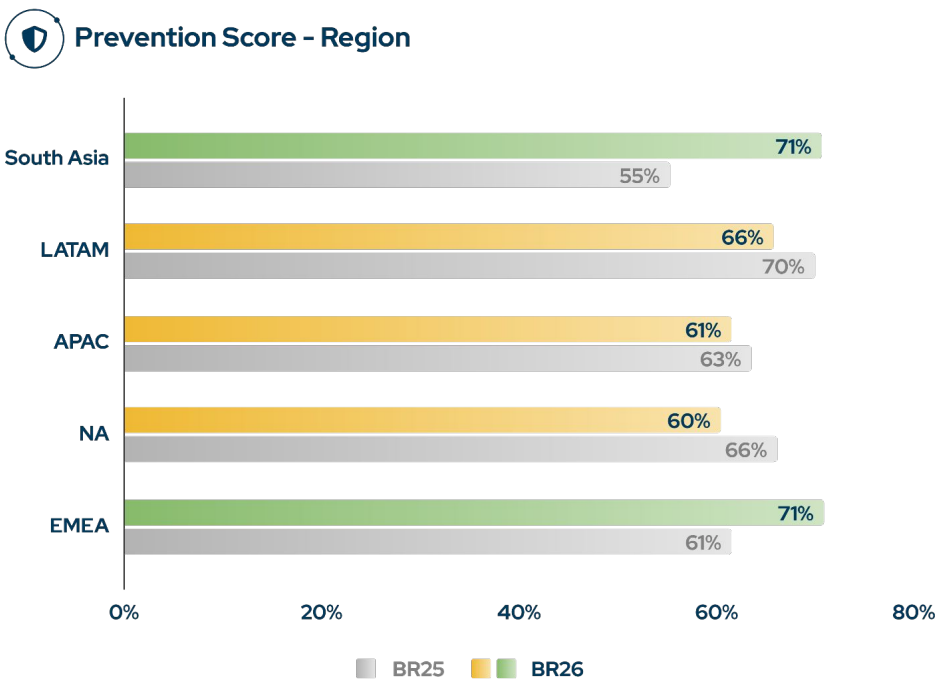
The pattern across industries is consistent with the global picture: log scores are rising in most sectors, sometimes dramatically, while alert scores cluster stubbornly between 13% and 33%. The log-to-alert conversion gap remains a systemic issue that no industry has solved.



Performance by Region

Prevention Effectiveness

In 2026, prevention effectiveness scores ranged from moderate to managed levels across regions, with notable shifts from the previous year and a striking change in leadership.



South Asia, the most exposed region in last year's report at 55%, **recorded the largest regional improvement of the year, rising 16 points to 71%** and moving from last place to a share of first. The gain suggests that the systemic control maintenance challenges flagged in the previous edition prompted meaningful investment in validation and remediation. **EMEA matched it at 71%**, up 10 points from 61%, recovering the leadership position it lost last year and reversing the resource-gap-driven decline observed in 2025.

Latin America (LATAM), last year's strongest performer at 70%, **slipped to 66%**, a reminder that regional leadership is as perishable as any other form of security advantage. **Asia-Pacific (APAC) eased from 63% to 61%**, continuing a pattern of solid but plateauing performance that points to controls in need of re-validation.

North America (NA) declined from 66% to 60% and now registers the lowest prevention effectiveness of any region. For a market dense with security investment and vendor presence, the result underlines a theme that runs throughout this report: spending on controls and validating controls are different activities, and only the second one shows up in prevention scores.

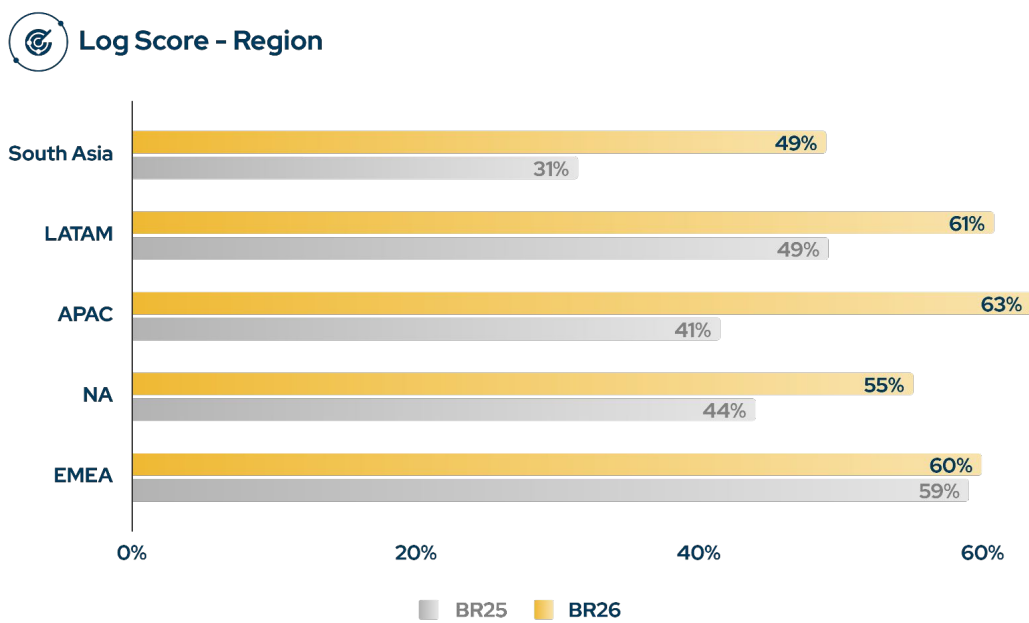
Across all regions, these results reinforce that control effectiveness is not static. Regions that trailed last year lead this year, and last year's leaders have slipped. Without frequent validation and tuning, even high-performing regions risk regression over time.



Detection Effectiveness

Detection performance in 2026 improved on the visibility side in every region, while alerting moved backward almost everywhere. The result is a global widening of the log-to-alert gap.

APAC led all regions in log score at 63%, a 22-point improvement from 41%, which represents the largest regional visibility gain of the year. However, its alert score declined from 16% to 13%, so the region's expanded telemetry has yet to translate into detection outcomes.

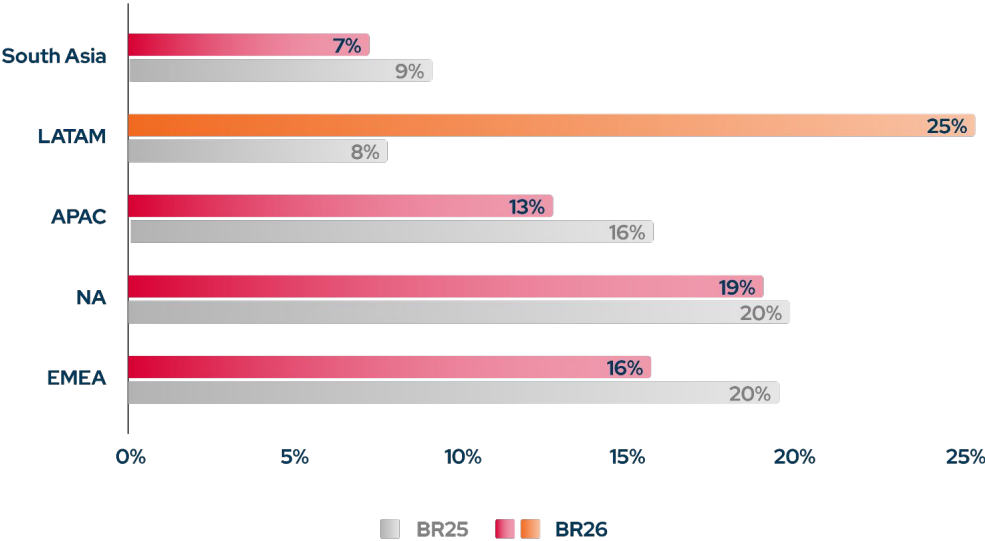


LATAM combined a strong log score improvement, from 49% to 61%, with the year's only significant alerting gain: its **alert score jumped from 8% to 25%**, the highest of any region. LATAM's progress across both prevention visibility and alerting last year and detection this year shows a region systematically maturing its pipeline, converting new telemetry into rules and alerts rather than letting it accumulate unused.

EMEA posted a log score of 60%, essentially flat against last year's region-leading 59%, but its alert score fell from 20% to 16%. **North America improved its log score from 44% to 55%**, yet its alert score slipped from 20% to 19%, leaving it second in alerting behind LATAM.



Alert Score - Region



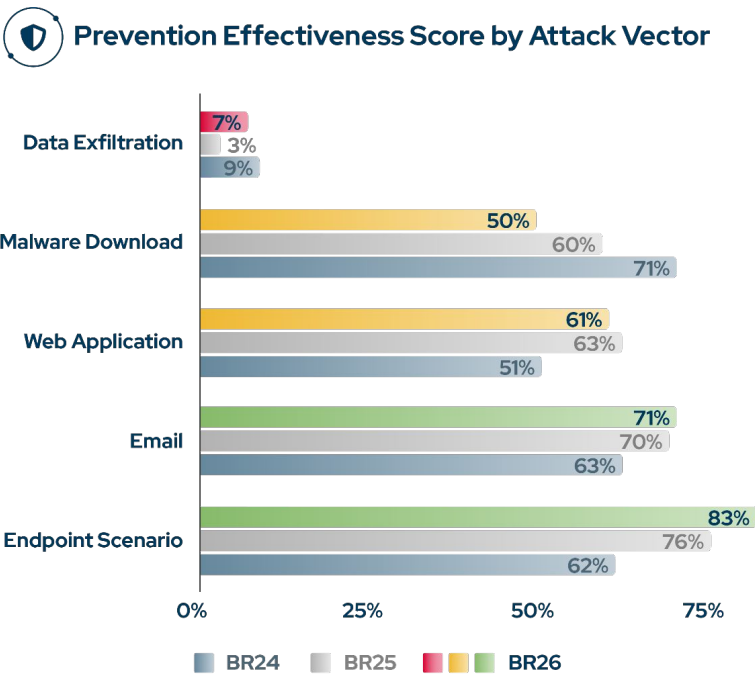
South Asia improved its log score substantially, from 31% to 49%, but it remains the lowest of all regions, and its alert score declined from 9% to 7%, again the lowest globally. Combined with its much-improved prevention score, the region presents an unusual profile: it now blocks attacks at a leading rate but remains largely blind to the ones that get through. Detection readiness is the region's clear next frontier.

While every region is progressing in telemetry coverage, detection engineering remains the weak point globally. Bridging the gap between log ingestion and actionable alerting requires not only better tools but continuous validation of rule logic, alert thresholds, and log source health tailored to each region's operational context.



Performance by Attack Vector

This year's findings continued to vary widely depending on the type of attack vector, highlighting the importance of validating controls across the entire kill chain. The overall prevention recovery was driven almost entirely by endpoint, email, and web defenses, while the vectors tied to data theft and payload delivery either stayed critically weak or got worse.



Once again, **data exfiltration emerged as the least prevented attack vector**, now for the fourth consecutive year. Its prevention score recovered to 7%, up from last year's low of 3%, though still below the 9% recorded in 2024. The improvement is welcome but must be kept in proportion: at 7%, more than nine out of ten simulated data theft attempts succeeded. With infostealer activity continuing to surge and ransomware operators relying on double extortion techniques to pressure victims, outbound data monitoring remains the most consequential gap in enterprise defense. Most organizations still lack the granular outbound monitoring and detection logic needed to block covert data theft.

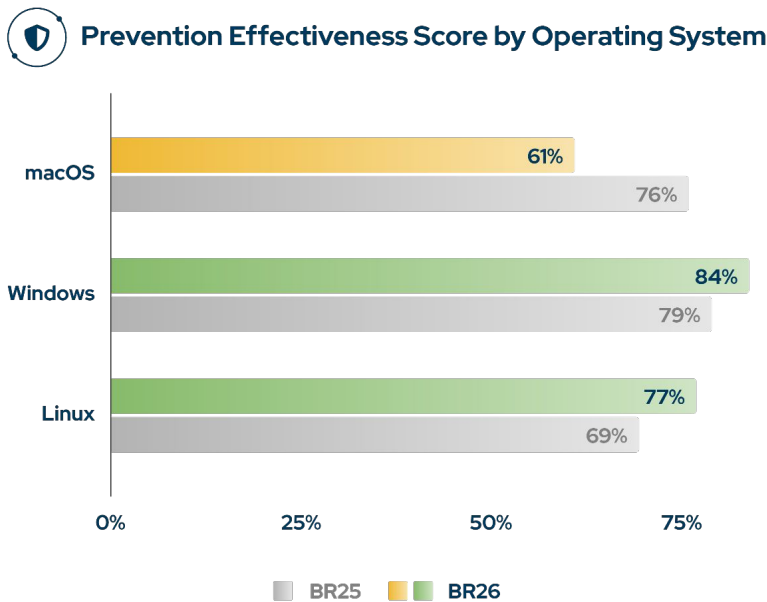
Prevention against endpoint scenarios improved further in 2026, reaching 83%, up from 76% last year and 62% the year before. Three consecutive years of gains show that sustained investment in EDR solutions, policy hardening, and post-compromise validation is paying off. Security teams are increasingly effective at identifying and stopping post-compromise activity, particularly around privilege escalation, a trend mirrored in this year's tactic-level results.

Email-borne threats remained one of the most targeted vectors in 2026, and prevention effectiveness edged up from 70% to 71%, holding the gains built through improved phishing detection, attachment sandboxing, and URL analysis capabilities across email security gateways. Continued testing of email modules and coverage of business email compromise (BEC) tactics remains essential.

Web application attacks were prevented 61% of the time, down slightly from last year's 63% though still well above the 51% recorded in 2024. The small decline suggests WAF tuning and input validation testing have plateaued, and given the widespread reliance on web platforms for customer-facing operations, even a modest prevention gap can result in substantial risk.



Prevention effectiveness for malware download scenarios declined to 50%, down from 60% last year and 71% in 2024. A 21-point slide over two years against one of the oldest and most well-understood attack vectors is this year's most troubling vector-level trend. As loaders and droppers continue to evolve, and as malicious payloads are increasingly delivered through legitimate-looking or compromised infrastructure, static content inspection is falling behind. Frequent validation of download scenarios remains critical to prevent initial access and downstream compromise, and this year's data shows that organizations that neglect it are losing ground at speed.

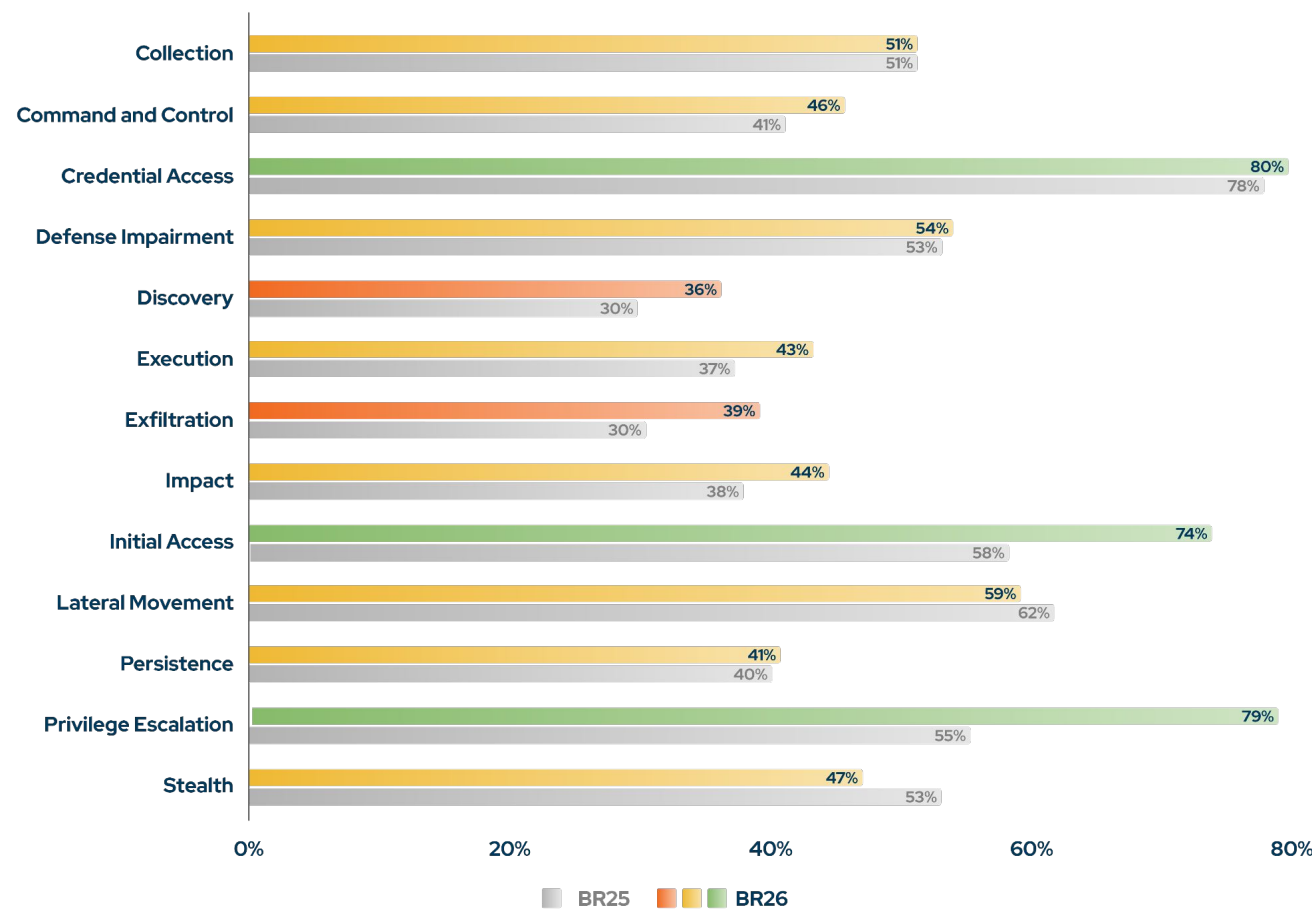


A breakdown by operating system, however, shows that the endpoint story is not uniform. **Windows endpoints led at 84%**, up from 79%, and **Linux systems improved from 69% to 77%**. macOS endpoints, by contrast, fell from 76% to 61%. Last year's report celebrated the dramatic recovery of macOS protection from 23% to 76%; this year, much of that gain has been surrendered, and macOS is once again the least protected endpoint platform. The lesson is uncomfortable but clear: the sector-wide investment that lifted Apple environments was not consolidated through continuous validation, and effectiveness decayed. Organizations with significant macOS estates should treat this regression as a call to re-test platform coverage now.



Performance by MITRE ATT&CK® Tactics

The MITRE ATT&CK® framework continues to serve as a strategic lens for assessing security readiness across the entire adversary kill chain. **In 2026, Picus simulations once again measured prevention effectiveness across the tactics of the ATT&CK Enterprise matrix.** The results show broad improvement across most of the kill chain, with dramatic gains in privilege and access defense, but they also confirm a stubborn divide: quiet, low-noise attacker behaviors remain far harder to stop than loud ones.





As in previous years, **Discovery was the least prevented tactic, with a prevention effectiveness score of 36%**. This is a six-point improvement from last year's 30%, but it still reflects the continued difficulty organizations face in stopping low-noise reconnaissance behaviors such as account enumeration, host discovery, and network scanning, activities that frequently go unnoticed until later stages of an attack.

- **Exfiltration followed at 39%**, up nine points from 30%. The improvement aligns with the modest recovery in the data exfiltration attack module and shows early progress in outbound monitoring, though the score remains far too low given the central role of data theft in modern extortion campaigns.
- **Persistence (41%, up from 40%), Execution (43%, up from 37%), and Impact (44%, up from 38%)** also ranked among the least prevented tactics. These stages involve actions such as script execution, payload deployment, service abuse, and sabotage, all of which remain difficult to catch without tightly configured endpoint controls and host-level monitoring. The gains are real but incremental, and the scores confirm that many organizations still struggle with visibility and control inside the endpoint layer after initial access.
- **Command and Control improved from 41% to 46%, and Collection held flat at 51%. Defense Impairment edged up from 53% to 54%. Stealth, by contrast, declined from 53% to 47%**, one of only two tactics to fall this year. The drop in Stealth prevention is significant: it indicates that evasion-oriented behaviors, from obfuscation to masquerading as legitimate activity, are improving faster than the behavioral analytics meant to catch them.
- **Lateral movement was the other declining tactic, easing from 62% to 59%**. After featuring as a relative strength in last year's report, the slip is a reminder that segmentation and identity controls need the same continuous re-validation as any other defense.

The strongest movements of the year came at the access-oriented end of the kill chain. **Initial Access jumped from 58% to 74%**, a 16-point gain that reflects the maturing of email security, web filtering, and exploit prevention observed in the attack vector data. **Privilege Escalation rose 24 points, from 55% to 79%**, the largest improvement of any tactic, consistent with the endpoint hardening and attack path remediation trends noted throughout this report.

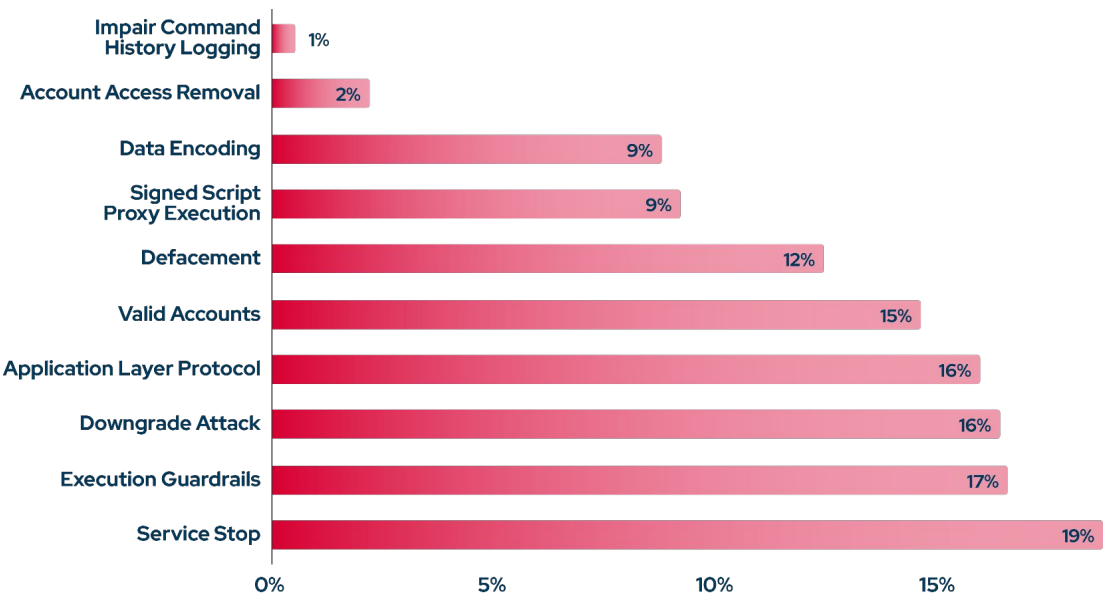
The standout performer remains Credential Access, at 80%, up from 78%. Password policy enforcement, credential hardening, and memory protection mechanisms continue to limit adversary access to authentication material. However, the persistently high success of downstream credential abuse, seen in last year's password cracking and Valid Accounts findings, indicates that what happens after a credential is obtained remains a serious concern.

In conclusion, the 2026 tactic results describe defenses that have become markedly better at stopping attackers from getting in and escalating, while remaining weakest against attackers who are already inside and staying quiet. Discovery, Exfiltration, Persistence, and Stealth form the modern adversary's safe operating space. To truly align with the ATT&CK framework, organizations must continuously validate not only control presence but control performance at every phase of the attack lifecycle, with particular attention to the low-noise behaviors where prevention still fails most often.



Performance by MITRE ATT&CK® Techniques

Beyond prevention at the tactic level, this year's technique-level findings expose where defenses fail most acutely against specific, well-documented adversary behaviors. The pattern is consistent with the tactic results: organizations are strong at the access-oriented front of the kill chain but weak against the quiet techniques used to evade logging, manipulate outbound traffic, and disrupt systems from inside. Every technique in this year's ten least prevented scored below 20%, meaning that for each one, the large majority of simulated attacks succeeded.



The single least prevented technique in 2026 was **Impair Command History Logging (T1562.003)**, blocked in just 1% of simulations. By clearing or disabling shell history, adversaries erase the record of their own commands before defenders or detection logic can capture it. That this ranks last is a direct, technique-level echo of the report's central detection finding: when the behavior that undermines logging is itself almost never prevented, the visibility gains reflected in the higher log score are easier for a capable attacker to neutralize than the aggregate number suggests.

Credential and Account Manipulation techniques followed close behind. **Account Access Removal (T1531)** was prevented only 2% of the time. Used to lock legitimate users out by deleting or disabling accounts, it appears frequently in the disruptive, extortion-driven stages of ransomware operations, and its near-total success rate shows how little coverage exists for adversary actions against identity infrastructure itself. **Valid Accounts (T1078)** was prevented 15% of the time. This remains low in absolute terms, and it reflects the persistent difficulty of distinguishing an attacker using stolen or weak credentials from a legitimate user. As campaigns lean on password spraying, MFA fatigue, and token theft, validating controls against credential-based access stays essential to interrupting lateral movement and privilege escalation.



Command and control techniques again ranked among the least prevented. **Data Encoding (T1132)**, used to obfuscate outbound communications, **was blocked in only 9% of simulations**, and Application Layer Protocol (T1071), which hides command and control inside common protocols such as HTTPS, DNS, or SMTP, in 16%. These low scores reflect the difficulty of separating malicious traffic from legitimate communication in environments that lack deep packet inspection, TLS inspection, or behavior-based anomaly detection. Outbound channels remain a reliable path for adversaries once inside.

Defense evasion techniques continue to defeat traditional prevention. **Signed Script Proxy Execution (T1216) was prevented in 9%** of simulations, exploiting trusted, signed system binaries to run malicious code under the cover of legitimacy. **Execution Guardrails (T1480.001) was blocked in 17%**, allowing adversaries to keep payloads dormant until they confirm they are running against a real target rather than an analysis environment. **Downgrade Attack (T1562.010), prevented in 16%**, forces systems onto weaker protocol or logging modes that reduce both protection and visibility. Together these show endpoint and inspection controls struggling to model the conditions under which evasive code actually executes, a weakness consistent with this year's decline in the Stealth tactic score from 53% to 47%.

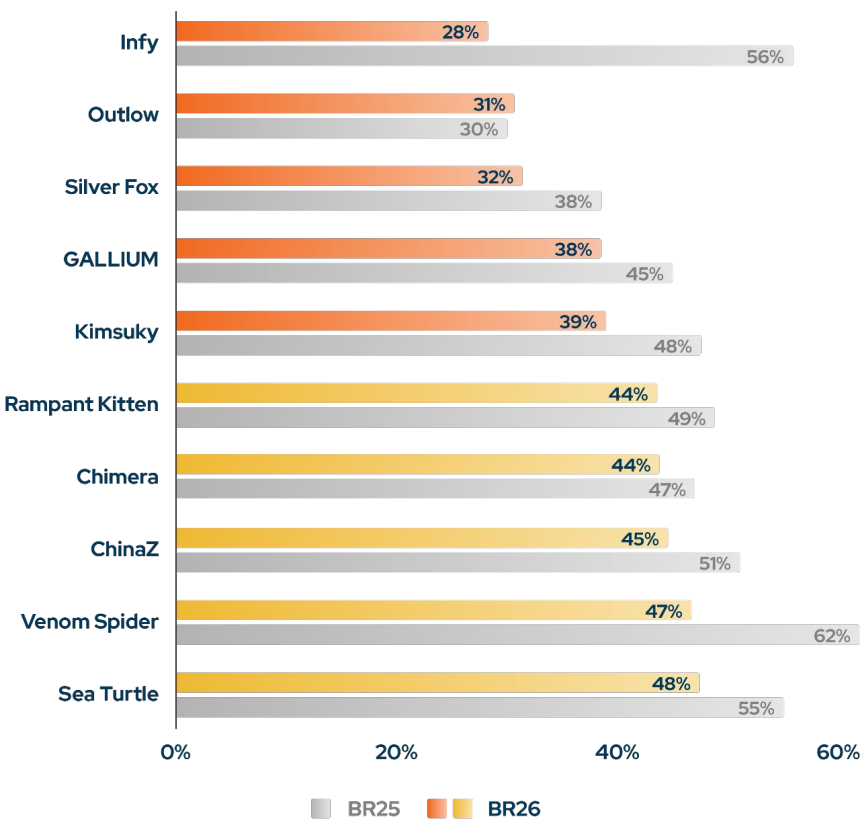
Impact techniques rounded out the list. **Defacement (T1491) was prevented in 12%** of simulations and Service Stop (T1489) in 19%. Both are used to disrupt operations and signal a successful breach, whether by altering public-facing content or halting the services and processes that keep systems and defenses running. Their low prevention rates indicate that many organizations still lack behavioral controls tuned to catch destructive and disruptive actions once an attacker has gained a foothold.

Taken together, the technique-level results reinforce the tactic-level story with sharper detail. The behaviors organizations prevent least are those that suppress logging, abuse valid identity, hide in normal traffic, and disrupt systems from within. These are precisely the low-noise, context-dependent actions that static, signature-led controls miss. Closing these gaps requires behavioral detection and continuous validation against the specific techniques adversaries rely on, rather than confidence built on control presence alone.



Performance by Threat Group

Organizations' ability to prevent real-world adversaries continues to vary significantly by threat group, and this year the variation moved in the wrong direction. Prevention effectiveness declined against nine of the ten least prevented groups, and the average score across the bottom ten fell from roughly 48% last year to roughly 40% this year. Every group in this year's bottom ten scored below 50%, meaning that for each of these adversaries, the majority of simulated attacks succeeded.



Infy, with a prevention score of just 28%, emerged as the most successful threat group in evading defenses this year. A long-running espionage actor associated with targeted surveillance campaigns, Infy was prevented at 56% in last year's analysis; the 28-point collapse is the sharpest deterioration of any group and shows how quickly defensive coverage erodes when an actor updates its tooling and delivery tradecraft faster than detection content evolves.

Outlaw, last year's least prevented group, remained near the bottom at 31%, essentially unchanged from 30%. Known for its opportunistic attacks across sectors, Outlaw blends cryptojacking, botnet activity, and lateral movement, often flying under the radar of conventional network controls. Two consecutive years at the bottom of the table indicate that its low-noise, commodity-infrastructure approach continues to defeat controls tuned for more conspicuous threats.

Silver Fox, an emerging APT group leveraging loader frameworks and obfuscated malware, followed at 32%, down from 38%. GALLIUM, an actor known for targeting telecommunications providers and their infrastructure, fell from 45% to 38%.



Established espionage groups also gained ground against defenders. **Kimsuky declined from 48% to 39%, Rampant Kitten from 49% to 44%, and Chimera from 47% to 44%.** These actors continue to exploit social engineering, macro-based payloads, and native system tools to avoid triggering traditional prevention logic. Their tactics emphasize quiet persistence and stealth, often bypassing controls focused only on known malware signatures, and this year's declining Stealth tactic score suggests exactly that dynamic at work.

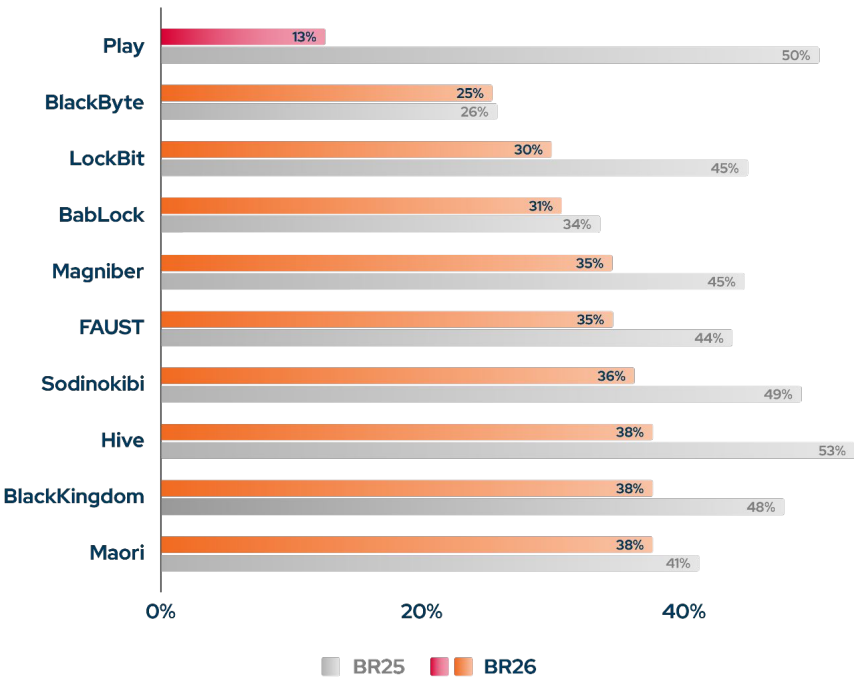
ChinaZ, a group historically associated with Linux-focused malware and DDoS botnets, **was prevented at 45%, down from 51%.** Venom Spider, known for its malware-as-a-service tooling used in targeted financial and e-crime campaigns, recorded one of the year's larger declines, falling from 62% to 47%. Sea Turtle, an espionage actor with a history of DNS-focused operations against government and telecommunications targets, rounded out the bottom ten at 48%, down from 55%.

The pattern across the bottom ten is consistent and concerning. These are not novel, unknown adversaries: most are well-documented groups with published tradecraft. Yet prevention effectiveness against them declined almost across the board, even in a year when overall prevention improved by seven points. The divergence shows that generic control improvement does not automatically translate into coverage against specific, adaptive adversaries. Organizations must validate their defenses against the threat groups most relevant to their sector and geography, using up-to-date emulations of each group's current behavior rather than last year's playbook.



Spotlight on Ransomware Attacks

Ransomware remains one of the most persistent and destructive threats facing organizations worldwide. In 2026, our research revealed that ransomware strains are evading defenses more effectively than at any point in this report's history. Every one of the ten least prevented ransomware families scored 38% or lower; the average prevention score across the bottom ten fell from roughly 44% last year to roughly 32% this year, and not a single family in the group showed meaningful improvement. Despite heightened awareness and continued investment in ransomware resilience, prevention effectiveness against leading ransomware families is moving backward.



Play ransomware is this year's most alarming result. Prevented at 50% in last year's analysis, Play collapsed to a prevention effectiveness of just 13%, making it the least prevented ransomware variant by a wide margin. The 37-point deterioration tracks the group's continued evolution in intermittent encryption, abuse of legitimate tools, and exploitation of public-facing applications. A strain that half of environments could stop a year ago now succeeds against nearly nine in ten.



BlackByte, the least prevented variant in each of the previous two editions, remained near the bottom at 25%, essentially unchanged from 26%. Known for exploiting public-facing applications and moving quickly to exfiltrate and encrypt sensitive data, BlackByte continues to bypass traditional controls, particularly in environments lacking behavioral monitoring and outbound traffic filtering. Three consecutive years at or near the bottom of this table make BlackByte the clearest standing indictment of static, signature-led ransomware defense.

LockBit, despite years of law enforcement attention and extensive public documentation, was prevented in only 30% of simulations, down sharply from 45%. BabLock followed at 31%, down from 34%, continuing to leverage double extortion techniques and staged execution patterns that evade static detection and sandbox inspection.

Magniber (35%, down from 45%) and FAUST (35%, down from 44%) showed similar prevention gaps, suggesting that many organizations still lack adequate coverage for file encryption behavior, registry modifications, and lateral movement within endpoint environments. Sodinokibi declined from 49% to 36%, a notable regression for one of the most analyzed ransomware families in existence.

Hive (38%, down from 53%), BlackKingdom (38%, down from 48%), and Maori (38%, down from 41%) completed the bottom ten. Hive's 15-point decline is particularly striking given the public takedown of its original operation and the volume of intelligence available about its tradecraft; its continued and growing success in simulation demonstrates that documentation is not defense.

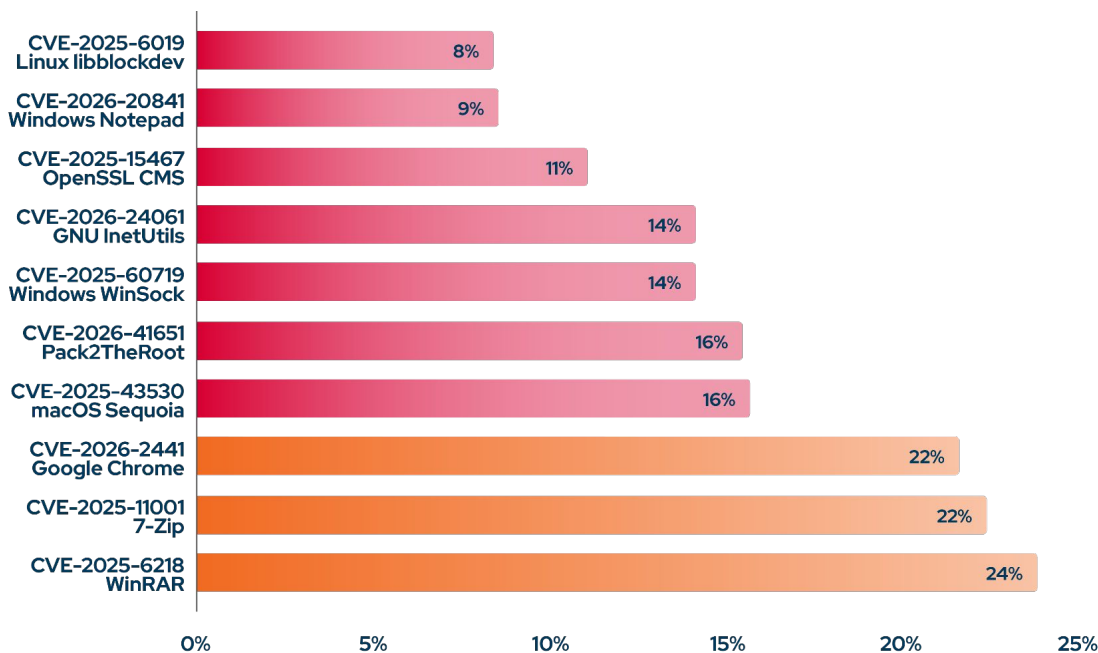
Taken together, these results are a clear warning. The broad improvement in endpoint scenario prevention recorded elsewhere in this report did not extend to the specific behavior chains that current ransomware families use. Organizations must simulate complete, current ransomware kill chains, including initial access, defense evasion, data exfiltration, and encryption or extortion stages, and validate that controls interrupt them at multiple points. Ransomware defense that was adequate against last year's variants is demonstrably failing against this year's.



Spotlight on Vulnerabilities

For the 2026 edition of the Blue Report, we continued to focus our vulnerability analysis on recently disclosed CVEs, covering vulnerabilities disclosed in 2025 and 2026. This approach offers a focused view into how well organizations are preventing exploitation of newly emerging threats and how quickly security controls adapt in the window between disclosure and weaponization.

The findings this year are worse than last year's. Every one of the ten least prevented vulnerabilities recorded a prevention effectiveness score below 25%, and the weakest were blocked in fewer than one in ten simulated exploit attempts. In last year's analysis, the upper end of the least-prevented list still reached the 55% to 59% range; this year, no vulnerability in the bottom ten comes close. The gap between disclosure and effective protection is widening, not closing, despite the availability of patches and detection signatures for most of these CVEs.



At the bottom of the list, the **Linux libblockdev CVE-2025-6019** vulnerability, a high-severity local privilege escalation issue (CVSS 7.0 High), was the least prevented vulnerability, with just 8% of simulated exploit attempts successfully blocked. Close behind was the **Windows Notepad CVE-2026-20841** vulnerability (CVSS 7.8 High) at 9%, a striking result for a flaw in one of the most ubiquitous components of the Windows operating system.

Two critical-severity vulnerabilities also ranked among the least prevented. The **OpenSSL CMS CVE-2025-15467** vulnerability (CVSS 9.8 High) was blocked in only 11% of simulations, and the **GNU InetUtils CVE-2026-24061** vulnerability (CVSS 9.8 High) in only 14%. Given the breadth of systems that depend on OpenSSL and core GNU network utilities, prevention rates this low against critical, network-reachable flaws represent substantial aggregate exposure.

The **Windows WinSock CVE-2025-60719** vulnerability (CVSS 7.0 High) was prevented 14% of the time, and the **Pack2TheRoot CVE-2026-41651** vulnerability (CVSS 8.8 High) 16% of the time. The **macOS Sequoia CVE-2025-43530** vulnerability, the only medium-severity entry on the list (CVSS 5.5 Medium), was blocked in just 16% of simulations, a result consistent with the broader macOS endpoint regression documented earlier in this report and a further signal that Apple environments are receiving less validation attention than they need.

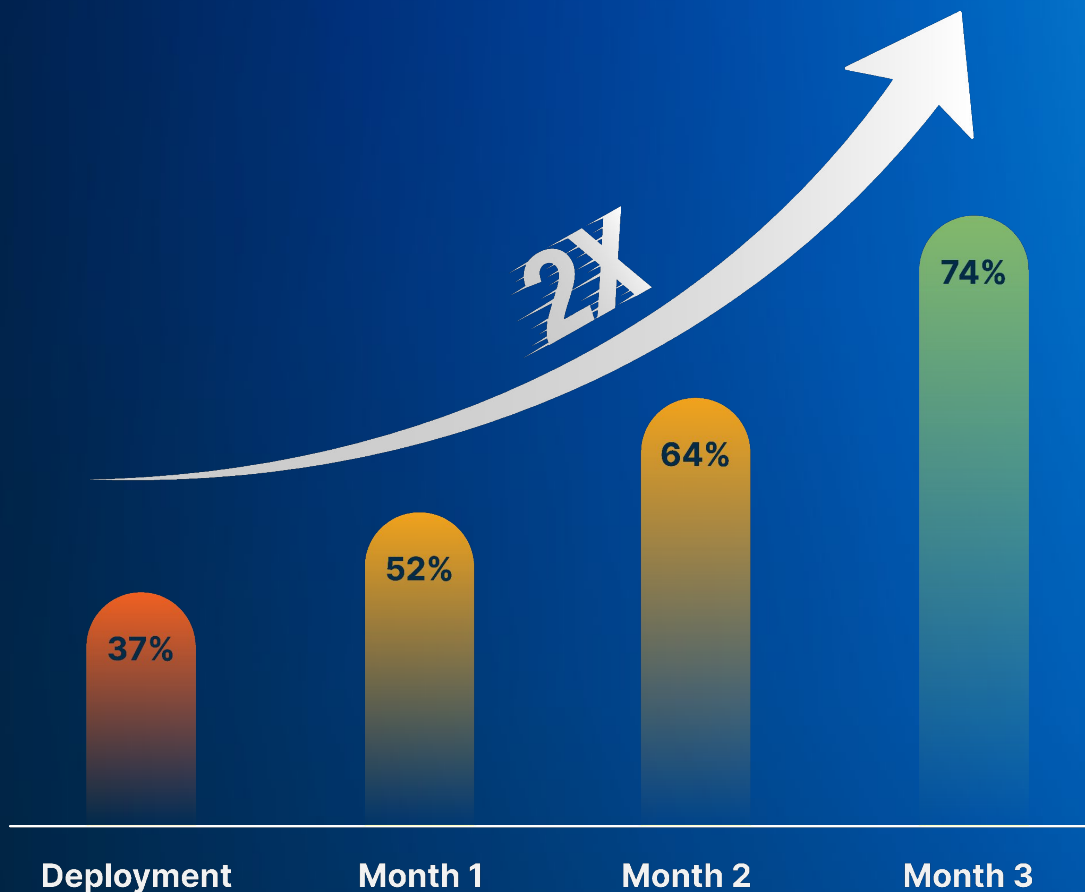


Even the best-prevented entries in the bottom ten remain deeply concerning. The **Google Chrome CVE-2026-2441** vulnerability (CVSS 8.8 High) and the **7-Zip CVE-2025-11001** vulnerability (CVSS 7.8 High) were each blocked in only 22% of simulations, and the **WinRAR CVE-2025-6218** vulnerability (CVSS 7.8 High) in 24%. These are among the most widely deployed applications in enterprise environments, with mature patching channels, and yet more than three-quarters of simulated exploit attempts against them succeeded.

The composition of this year's list tells its own story: a browser, two archive utilities, a text editor, core operating system components across Windows, Linux, and macOS, and foundational cryptographic and network libraries. These are not exotic enterprise appliances; they are the everyday software surface of virtually every organization. The consistently low scores suggest that many organizations remain slow to patch operating systems and common desktop software, and that compensating controls are not picking up the slack.

Whether through patching or mitigation, organizations must ensure that vulnerabilities are addressed before they can be weaponized. Where timely patching is not possible, it becomes essential to have effective compensating controls, such as endpoint protection, intrusion prevention systems, and application-layer filtering, validated against exploit attempts tied to recent CVEs. This is particularly important for zero-day and n-day vulnerabilities, where the window between public disclosure and exploitation is often too short for traditional patch cycles to keep up.

Picus Security Customers Prevent Twice As Many Attacks



Picus Security provides an autonomous exposure validation platform, the Picus Platform, powered by **Picus Breach and Attack Simulation**, **Picus Autonomous Penetration Testing**, and **Picus Exposure Validation**, with **Numi AI** and **Picus Swarm** as the intelligence and orchestration layer, to help organizations of all sizes continuously validate and reduce their cyber risk. Security teams can prove what attackers can actually exploit, evaluate the effectiveness of their security controls, discover at-risk assets, and identify high-risk attack paths that attackers could use to reach critical systems and users.

On average, our customers prevent twice as many attacks within just three months. With Picus, security leaders can quickly mature their security posture and move beyond basic vulnerability management. Instead of spending their days making impossible trade-offs that may leave gaps in their defenses, they can consistently and successfully defend against sophisticated multi-pronged attacks.

About Picus Security

Picus Security is the pioneer of Breach and Attack Simulation and is now defining Adversarial Exposure Validation, enabling organizations to understand and reduce cyber risk with precision and speed. The Picus Platform empowers security teams to continuously correlate, prioritize, and validate exposures across on-premises, cloud, and hybrid environments, autonomously and at machine speed.

By simulating real-world attacker behaviors and proving what defenses actually stop, Picus helps teams focus on the exposures attackers can truly exploit and the high-impact fixes that close them, rather than managing siloed alerts or theoretical risks. With ready-to-deploy mitigations and actionable guidance, the Picus Platform makes it easier to stop more threats with less effort, while a named human stays on every risk-acceptance decision.

For more information, visit picussecurity.com

BLUE REPORT™
2026

PICUS