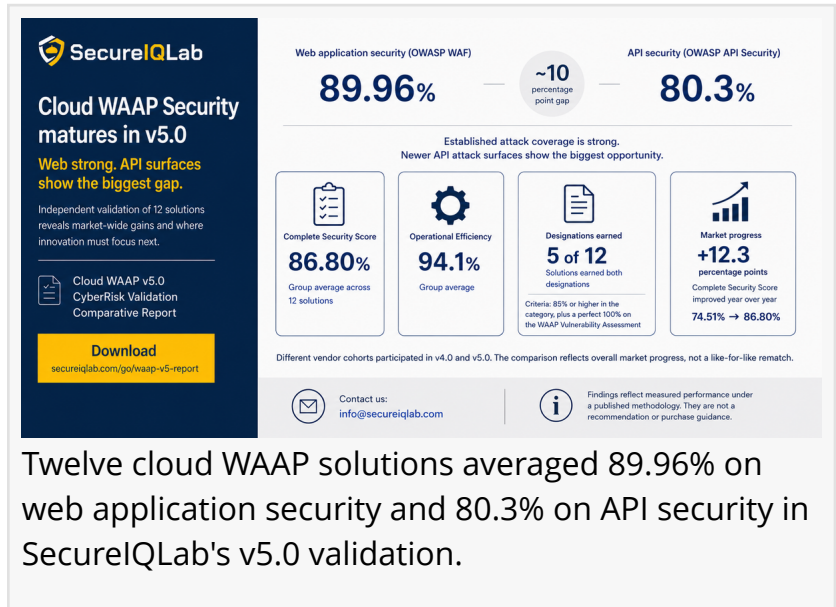


SecureQLab Report: Cloud WAAP Solutions Near Ceiling on Legacy Attacks, Lag on New API Surfaces

Complete Security Scores averaged 86.80 percent, up 12.3 percentage points year over year, but API security trailed web security by nearly 10 points.

AUSTIN, TX, UNITED STATES, August 3, 2026 /EINPresswire.com/ -- The [Cloud WAAP v5.0 CyberRisk Validation Comparative Report](#) from SecureQLab reveals a two-speed industry. Across 12 cloud Web Application and API Protection (WAAP) solutions, the group average Complete Security Score reached 86.80 percent, an increase of approximately 12.3 percentage points over the group average recorded in the v4.0 validation. The report notes that the vendor cohorts differed between the two evaluations, so the comparison indicates overall market progress rather than a like-for-like measurement of the same products.



Twelve cloud WAAP solutions averaged 89.96% on web application security and 80.3% on API security in SecureQLab's v5.0 validation.

“

The industry has mastered the attacks it has spent a decade testing. The gaps this validation measured sit exactly where attackers are moving next.”

David Ellis, VP of Research and Corporate Relations, SecureQLab

That headline number rests on uneven ground. The report puts the group at 89.96 percent on OWASP Web Application Security and 80.3 percent on OWASP API Security, close to a 10 percentage point gap between the surfaces enterprises have tested longest and the ones they are adding fastest. The report's own conclusion names Compliance and OWASP API Security as the two capabilities with the greatest variability across the 12 solutions.

The report replaces vendor-rated specifications with measured performance. Every product faced an identical evidence base: 1,608 attack payloads and 1,487 benign requests executed under controlled conditions, scored against the standards-

governed Cloud WAAP CyberRisk Validation Methodology v5.0, which maps outcomes to MITRE ATT&CK, OWASP Top 10 (2025), OWASP API Security Top 10 (2023), OWASP LLM Top 10, and AMTSO.

Key facts:

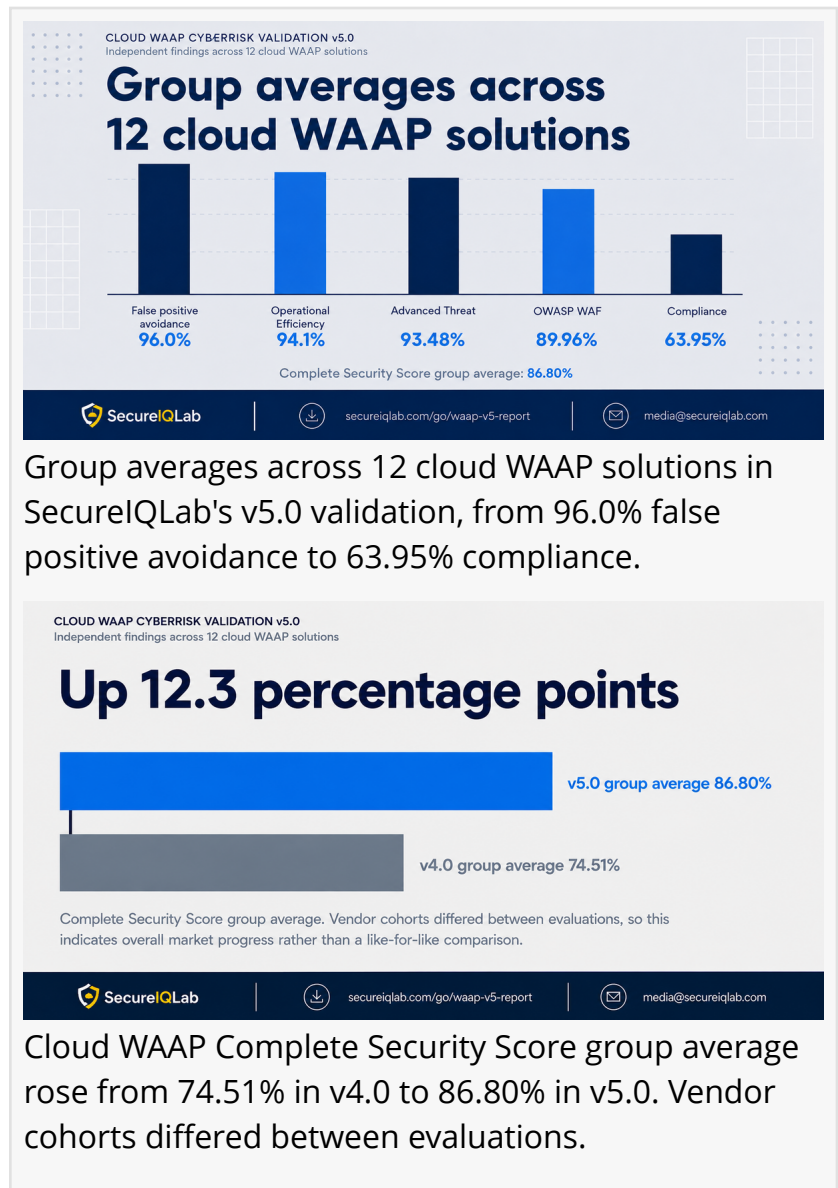
- 12 vendors engaged, each named with its full results in the report
- Security Efficacy covers 7 tested areas: OWASP Web Application Security, OWASP API Security, and five advanced threat categories (Bot Attacks, AI-Assisted Bot Attacks, Layer 7 DoS and DDoS Attacks, Security Resiliency, and WAAP Vulnerability Assessment)
- 6 validation pillars: Security Efficacy, Operational Efficiency, AI Application Security & Operational Efficiency, Secure-by-Design & Secure-by-Default, False Positive Avoidance, and Compliance
- First WAAP methodology to include large language model (LLM) security testing, and the first to use AI-enhanced attack payloads against AI-defended products
- AMTSO-compliant (Testing Protocol Standard v1.3), AMTSO Test ID: AMTSO-LS1-TP169

Where has the industry converged?

Legacy attack surfaces are close to solved. SecureQLab's category-level scoring puts the group at 97.9 percent on Layer 7 DoS attacks, with 11 of the 12 products at 100 percent, and at 99.7 percent on the WAAP Vulnerability Assessment. Authentication-failure detection under OWASP Top 10 (2025) averaged 99.5 percent. On the attack classes the industry has tested longest, performance is convergent and near the ceiling.

Where is the industry exposed?

The newer API attack surfaces show a different market. SecureQLab's category-level scoring records six of the 12 products at zero on WebSocket API protection, a group that includes at least



one product that does not support the protocol at all as well as products that support it and did not defend it. Privilege-escalation attempts against APIs, the Broken Object Level Authorization category, split the field. Six of the 12 products scored 100 percent and four scored zero, leaving that category average at 58.3 percent.

On the AI surfaces the picture is narrower than the range suggests. Prompt Injection (LLM01:2025) defense scores ran from 100 percent down to 35 percent, with nine of the 12 products at 100 percent, while Improper Output Handling (LLM05:2025) averaged 99.4 percent. The gap is technique-specific and does not extend to AI defense generally. Under AI-assisted bot attacks, most of the field held at 100 percent, and two products with perfect scores against traditional bot attacks dropped below full marks. The OWASP LLM security results are scored independently and are not factored into the Complete Security Score or the overall Operational Efficiency score.



SecureIQLab is an independent cloud security validation laboratory based in Austin, Texas. Unlike traditional analyst firms that rely on subjective surveys, SecureIQLab provides empirical, real-time security metrics based on testing that maps real-world e

The sharper AI finding is operational rather than defensive. The same solutions that blocked AI-targeted attacks nearly perfectly diverge widely on readiness to run AI security, with the report's separate Gen AI and LLM Operational Efficiency results ranging from 100 percent down to zero. For enterprises adding AI features to existing applications, the deployment and governance gap is wider than the detection gap.

What does the false-positive spread mean for operations?

Accuracy varies as widely as coverage. The false-positive evaluation sent every product the same 1,452 benign requests, with the 35 LLM benign cases scored separately. Four of the 12 products produced zero false positives, the group median was 9, and the highest count was 275, an 18.9 percent false-positive rate on the identical benign traffic. For security operations teams, that spread is a direct measure of alert-triage workload and blocked legitimate transactions.

The report also includes Compliance: regulatory, audit, logging, and governance validation across

web and API workloads. Compliance scores spanned 47.8 percent to 87.0 percent across the group and averaged 63.95 percent, with six of the 12 solutions above that average. SecureQLab's layer-level scoring puts the regulatory and AI security layers lowest. Operational Efficiency remained the industry's strongest result, with a group average of 94.1 percent across the nine evaluated Operational Efficiency categories and 10 of the 12 products above 90 percent. That aggregate conceals category-level spread the report publishes: every solution scored above 90 percent on Support and Documentation and on Visibility and Analytics, while Ease of Deployment ranged from 66.7 percent to 100 percent and Ease of Management from 72.7 percent to 100 percent.

Which products earned the Secure by Design and Secure by Default designations?

The report awards a Secure by Design badge to products that scored 85 percent or higher on the Secure by Design criteria and a perfect 100 percent on the WAAP Vulnerability Assessment, and it applies the same two thresholds to Secure by Default. Secure by Design was validated against three core principles and 22 criteria, Secure by Default against the first two principles and 13 criteria. Five of the 12 products earned both badges. The report names them.

"The industry has mastered the attacks it has spent a decade testing. The gaps this validation measured sit exactly where attackers are moving next. Those surfaces are API session logic, WebSocket channels, and prompt injection. Enterprises deserve measured evidence of protection for them rather than datasheet claims," said David Ellis, VP of Research and Corporate Relations at SecureQLab.

The full Cloud WAAP v5.0 CyberRisk Validation Comparative Report is available at secureqlab.com/go/waap-v5-report. It publishes each product's overall Security Efficacy and Operational Efficiency scores with its Cloud WAAP CyberRisk Ripple tier, results across the nine Operational Efficiency categories, Matthews correlation coefficient, precision and recall for each product, category-level group averages and score ranges, false-positive avoidance rates, and the Compliance results. Detail below that level will be published in the individual vendor reports.

Frequently Asked Questions

Q: How was the research conducted? A: SecureQLab executed 1,608 attack payloads and 1,487 benign requests against each of the 12 WAAP solutions under identical, controlled lab conditions, following the AMTSO-compliant Cloud WAAP CyberRisk Validation Methodology v5.0 (AMTSO Test ID: AMTSO-LS1-TP169).

Q: What is the most significant finding? A: SecureQLab's category-level scoring shows the performance gap between mature and new attack surfaces: group averages of 97.9 percent on Layer 7 DoS and 99.7 percent on the WAAP Vulnerability Assessment versus 48.3 percent on WebSocket API protection and 58.3 percent on API privilege escalation.

Q: Does the report rank vendors? A: Yes. All 12 solutions are assigned to one of four ranking tiers: Leader, Contender, Visionary, or Upcomer. Placement is derived from each solution's total Security Efficacy and Operational Efficiency scores and is plotted in the Cloud WAAP CyberRisk Ripple. Solutions are listed alphabetically within each tier.

Q: Is a tier assignment an endorsement? A: No. A tier records where measured scores fell under a uniform methodology applied identically to every product. It is not a recommendation to buy. Results are determined solely by measured performance under the published methodology. The per-category data exists so enterprises can weigh the surfaces that matter in their own environments against their own evaluation criteria.

Q: What is new in the WAAP 5.0 methodology? A: WAAP 5.0 introduced 3 brand-new evaluation areas (AI-assisted bot attacks, API gateway operational efficiency, LLM/GenAI security) and used AI-enhanced payload qualities for more sophisticated attack simulation.

Q: What should security teams do with these findings? A: Inventory which of the newer attack surfaces (WebSocket APIs, API authorization logic, LLM integrations) exist in their environments, then validate deployed protections against measured evidence for those specific surfaces.

About SecureQLab

SecureQLab is an independent cloud security validation laboratory based in Austin, Texas. Unlike traditional analyst firms that rely on subjective surveys, SecureQLab provides empirical, real-time security metrics based on testing that maps real-world enterprise use cases to specific business challenges. SecureQLab is a principal member of Mplify (formerly MEF) and a member of the Anti-Malware Testing Standards Organization (AMTSO), AVAR, and NetSecOPEN.

Data Integrity Disclosure: SecureQLab does not endorse specific vendors. The findings in this report represent objective data captured under controlled conditions using the published validation methodology. These results are presented as verified performance metrics and do not constitute purchase guidance for any product. Ranking-tier placement records where a product's measured scores fell under the published methodology and is not a recommendation to buy. SecureQLab disclaims all warranties regarding the application of this data to unique user environments.

SecureQLab Communications

SecureQLab

+1 512-575-3457

[email us here](#)

Visit us on social media:

[LinkedIn](#)

[Bluesky](#)

[Instagram](#)

This press release can be viewed online at: <https://www.einpresswire.com/article/931138099>

EIN Presswire's priority is source transparency. We do not allow opaque clients, and our editors try to be careful about weeding out false and misleading content. As a user, if you see something we have missed, please do bring it to our attention. Your help is welcome. EIN Presswire, Everyone's Internet News Presswire™, tries to define some of the boundaries that are reasonable in today's world. Please see our Editorial Guidelines for more information.

© 1995-2026 Newsmatics Inc. All Right Reserved.