



Continuous Adversarial Security

Not a point-in-time test, an always-on adversary working for you.

SERVICE PACKAGE

August 2026

fletchlabs.ca • info@fletchlabs.ca • Cambridge, Ontario, Canada

Executive summary

The annual penetration test is a snapshot; your environment is a moving picture. Continuous Adversarial Security replaces the once-a-year assessment with an ongoing adversarial cycle of Recon, Discover, Validate, Exploit, Report, and Retest, run continuously against your environment by a named Fletch Labs operator who learns it more deeply with every pass. The goal is simple: find and fix exploitable weaknesses before an adversary does, instead of waiting for detection tooling to catch an intrusion already in progress.

Why detection alone is a losing battle

Managed detection and response (MDR), XDR, EDR, and SIEM are necessary layers, and we deploy and tune them for clients every week. But they are fundamentally reactive: they observe an attack that is already underway. The data shows that attacker tempo has outrun the patch-and-detect model, and generative AI is compressing those timelines further.

63 → 5 days

Mandiant's measured average time-to-exploit for newly disclosed vulnerabilities collapsed from 63 days (2018–2019) to just five days in 2023. Attackers now move quickly enough to beat routine patching cycles.³

#1 vector

Exploits, both zero-day and n-day, were the number-one initial infection vector in Mandiant incident-response engagements for four consecutive years (2020–2023). Compromise increasingly starts at an unassessed weakness, not a phished user.³

Top entry

The 2026 Verizon Data Breach Investigations Report finds software-vulnerability exploitation has overtaken stolen credentials as the leading way attackers get in, and documents dozens of attack techniques now bolstered by generative AI, from spotting security gaps to writing malware.¹

\$4.99M

The global average cost of a data breach hit a record USD \$4.99 million in 2026, up 12% year over year, driven by higher detection, escalation, and lost-business costs. The same research recorded a 56% increase in AI-driven attacks, led by deepfake impersonation and AI-enabled malware.²

Nation-state actors are already using large language models for reconnaissance, vulnerability research, scripting, and social engineering.⁴ As offensive AI lowers the cost of finding and weaponizing flaws, the defensive question changes from *will my tooling catch the intrusion?* to *who finds my exploitable weaknesses first: my security partner, or the adversary?* Proactively and continuously assessing your own environment is the highest-leverage way to stay ahead of that curve.

The service: a continuous adversarial cycle

Continuous Adversarial Security runs the following loop against your environment on an ongoing basis. Each phase feeds the next, and each full cycle feeds the one after it:

- **Recon:** Continuous asset and attack-surface discovery, including shadow IT and drift from the prior cycle, so each loop starts from where your environment actually is today.
- **Discover:** Vulnerability and weakness identification across the current surface, surfaced as it changes rather than once a year.
- **Validate:** Manual confirmation of every finding before it reaches your team, eliminating false positives and noise up front.
- **Exploit:** Controlled exploitation to prove real-world impact, with formal penetration testing applied only where explicitly scoped and authorized.
- **Report:** Prioritized findings tied to business impact, not only CVSS scores, so remediation reflects what an actual attacker would do.
- **Retest:** Verification that remediations genuinely closed the gap, closing the loop and feeding the next cycle with what we learned.

Because the loop never stops, the value compounds. We track what changes between cycles, surface exploitable risk as it emerges, and verify that fixes actually closed the gap, rather than handing you a static report that is stale the day a new service ships.

Real operators, not autonomous agents

Every engagement is led by a **named Fletch Labs representative** who works directly with your team, remains your point of contact across cycles, and carries forward the accumulated knowledge of your environment. We conduct actual operations: every exploitation step is scoped, authorized, and executed by an accountable human operator. We do not point an autonomous AI agent at your network and let it run. Tooling assists, people decide. That is both a safety property and a quality property: validated findings, zero noise, and judgment applied where automation cannot be trusted.

Engagement & pricing

Delivered as a flat-rate subscription, billed monthly. One rate covers a single primary environment under standard scope, with no per-host or per-user math. Environments significantly larger than a typical small-to-mid-size attack surface may need a brief scope conversation before signing, so the engagement stays sustainable on both sides.

Term	Monthly rate	Total contract value
6-month term	CAD \$2,750 / month	CAD \$16,500
12-month term	CAD \$2,500 / month	CAD \$30,000

Cancellation

The initial term (6 or 12 months, as selected) is a committed period. The terms below cover what happens at the end of that term and how cancellation works afterward.

- The initial term is not cancellable for convenience; early exit is available only for cause (material breach) or by mutual agreement.
- At the end of the initial term, the contract automatically moves to month-to-month at the 6-month rate unless renewed into a new term.
- Once month-to-month, either party may cancel with 30 days' written notice.
- No refunds for partial months already in progress.

Configurable scope: autonomous vs. coordinated

Not every organization is comfortable with fully autonomous exploitation running against production, and that is expected. Scope is configured up front and can be adjusted at renewal:

- **Autonomous mode** runs recon through exploitation continuously without per-action sign-off, with results delivered as they are found for teams that want maximum coverage and minimum friction
- **Gated mode** runs recon, discovery, and validation continuously, but exploitation requires a go / no-go check-in before anything is attempted against production for environments where an untimed exploit attempt could cause business disruption
- **Coordinated mode** runs the full loop on a scheduled cadence agreed with your team (e.g. a defined testing window), functioning closer to a formal, pre-approved penetration test repeated on a regular cycle than a fully autonomous background process for teams that prefer a pre-approved testing window on a recurring schedule

Start a conversation

Tell us about your environment and current testing cadence, including which scope mode fits your team, and we will map out what a continuous adversarial cycle would look like for you.

Reach us at info@fletchlabs.ca or through fletchlabs.ca/contact.

Sources

1. Verizon, 2026 Data Breach Investigations Report — <https://www.verizon.com/business/resources/reports/dbir/>
2. IBM & Ponemon Institute, Cost of a Data Breach Report 2026 — <https://www.ibm.com/reports/data-breach>
3. Mandiant (Google Cloud), How Low Can You Go? An Analysis of 2023 Time-to-Exploit Trends — <https://cloud.google.com/blog/topics/threat-intelligence/time-to-exploit-trends-2023>
4. Microsoft Threat Intelligence & OpenAI, Staying ahead of threat actors in the age of AI — <https://www.microsoft.com/en-us/security/blog/2024/02/14/staying-ahead-of-threat-actors-in-the-age-of-ai/>