

Why Isolation Forest: Behavioral Machine Learning Detection Explained

The Algorithm, the Two-Feature Design, and Why It Outperforms Signatures and Rules for This Threat

Bobby Smathers - SimSpace Principal Security Architect

July 2026

1. Purpose

This document is a deeper technical companion to the architecture whitepaper's detection-engine section, written to stand on its own as a reference or appendix. It covers three things in detail: how Isolation Forest actually works mechanically, why a two-feature design (latency and response size alone) is a deliberate and correct engineering choice rather than a simplification, and why this behavioral, active-probing approach is fundamentally better suited to this threat than signature- or rule-based detection, while being honest about where those older approaches still earn their place.

2. What Isolation Forest Actually Does

Isolation Forest is built on a single, almost deceptively simple observation: anomalies are, by definition, few and different. If you repeatedly split a dataset in half using random cuts along random features, a point that sits in a dense, ordinary cluster takes many splits to separate from its neighbors, because there's always another nearby point to keep sharing a partition with. A point that sits far from everything else gets isolated almost immediately, often in just a handful of random cuts, because there's no crowd around it to keep it company.

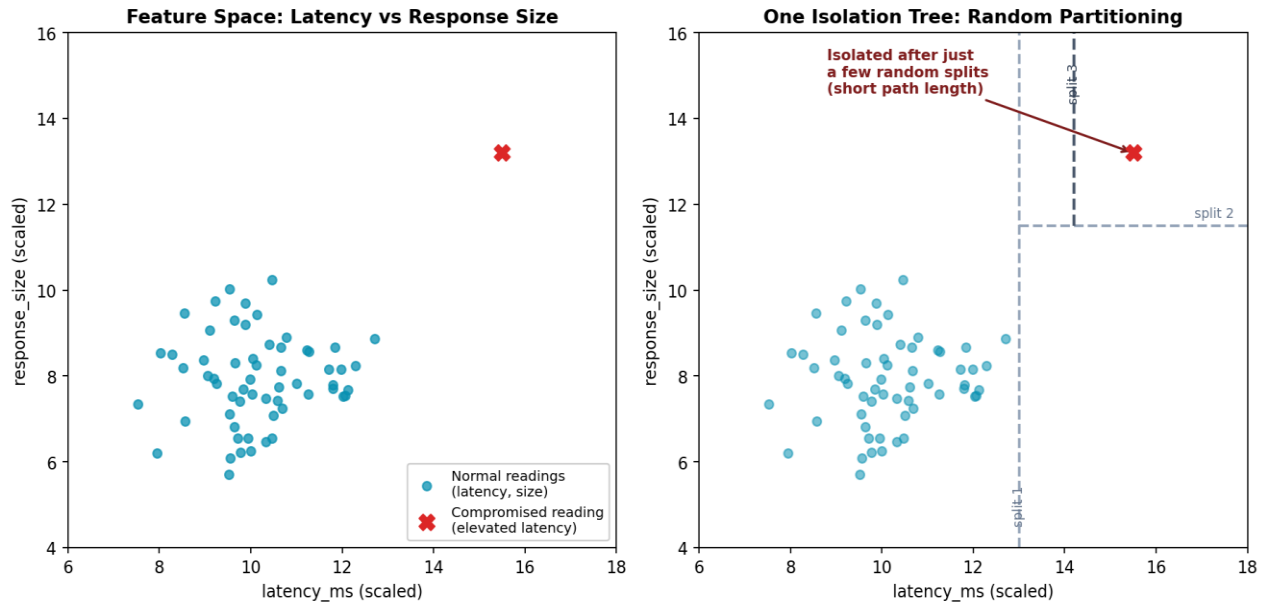


Figure 1. Left: a compromised reading sits far from the normal cluster in latency/size space. Right: a single isolation tree separates it from everything else in just three random splits, while an ordinary point nearby would take many more.

An isolation tree formalizes this: at each node, pick a random feature and a random split value within that feature's current range, and partition the data. Repeat recursively until every point is alone in its own leaf. The number of splits it took to isolate a given point is its path length in that tree. Build an ensemble of many such trees (100, in this implementation), each with different random splits, and average the path length across all of them. Short average path length means the point isolates easily and is unusual; long average path length means it's buried in a crowd and is ordinary.

That averaged path length converts to a bounded anomaly score via the field's standard formulation:

$$s(x, n) = 2^{-(E(h(x)) / c(n))}$$

where $E(h(x))$ is the average path length across the ensemble for point x , and $c(n)$ is an expected-path-length normalization term based on the training set size n , borrowed from the average search-path length of an unsuccessful search in a binary search tree. Scores approach 1 for clear anomalies, approach 0.5 or below for points that are entirely unremarkable, and the calibration is done once, at fit time, purely from the shape of the baseline data itself.

2.1 Why This Doesn't Need Labeled Attack Data

This is the property that matters most for a threat like this one. Isolation Forest is unsupervised: it never needs an example of what an attack looks like. It only needs a sample of normal operation. It learns the shape of ordinary behavior and flags whatever doesn't fit that shape, whether or not anyone has ever seen that specific kind of deviation before. A supervised classifier, by contrast, can only recognize categories of attack it was explicitly

trained on; anything sufficiently novel is invisible to it by construction. For a threat where the whole premise is that new CVEs and new exploitation techniques appear continuously, that distinction is the entire ballgame.

2.2 Why an Ensemble, Not a Single Tree

A single isolation tree's random splits could, by chance, happen to isolate an ordinary point quickly or bury a genuine anomaly under an unlucky sequence of cuts. Averaging over 100 independently randomized trees smooths out that variance; a real anomaly isolates quickly across most of the ensemble, not just one lucky tree, while an ordinary point that got unlucky in one tree gets corrected by the other 99. This is the same bias-variance logic behind any ensemble method, applied here to isolation depth instead of a class vote.

3. Why Isolation Forest, Specifically

Isolation Forest is not the only unsupervised anomaly detection method available. It was chosen deliberately over the alternatives, for reasons specific to this problem.

Approach	Why it's a worse fit here
One-Class SVM	Training cost scales poorly with data volume (effectively quadratic or worse), and results are sensitive to kernel and parameter choice in ways that are hard to explain to a training audience or re-tune confidently on new hardware.
Density-based (LOF, DBSCAN)	Requires computing local density/distance relationships between points, which gets expensive as data grows and struggles when normal behavior itself has varying density (bursty legitimate traffic against a quiet baseline, exactly this project's traffic pattern).
Autoencoders / deep learning	Needs substantially more training data to learn anything meaningful, is materially harder to explain per-decision (no equivalent of tree votes or path depth), and is disproportionate machinery for a two-feature problem.
Simple z-score / statistical threshold	Only evaluates one feature at a time and assumes a roughly Gaussian distribution; can't naturally express a joint relationship between latency and response size, and is brittle against real-world traffic that isn't cleanly bell curved.
Isolation Forest	Linear-time training ($O(n \log n)$), no distance or density computation, naturally handles multiple features jointly, and produces a genuinely explainable per-reading result: which trees flagged it, how deep, and how that compares to the baseline.

That last point, explainability, is not a minor bonus. A detector built for a training exercise needs to be able to answer "why did this alert fire" in terms a human operator can actually follow, in real time, without opening a debugger. Isolation Forest's tree-vote and path-depth structure gives that answer directly; several of the black-box alternatives above cannot.

4. Why Two Features Is Correct, Not a Shortcut

The engine scores exactly two features: request latency and response payload size. Earlier iterations carried five, including HTTP status and jitter. Cutting three of them was not simplification for its own sake; each one was removed for a specific, testable reason.

- HTTP status was a constant 200 on every safe-path probe. A feature with zero variance carries zero information; including it does nothing but add dead weight to the feature vector.
- latency_deviation (the absolute difference from an assumed 30ms baseline) is a pure linear function of latency itself above that point. Feeding the model both latency and a linear transform of latency is the same signal counted twice, not two independent signals; redundant features dilute a distance- or split-based model rather than strengthening it.
- jitter_ms (time since the last poll of that device) was cut after directly causing a false-positive flood in testing. Under a steady polling cadence, jitter has an extremely tight natural spread, so ordinary scheduling noise easily registers as an extreme deviation on that axis alone, with no causal connection to compromise. A device under attack shows elevated processing latency, not irregular polling; jitter was measuring something the threat model doesn't actually produce.

The general principle, not specific to this project: more features are not automatically better in an unsupervised model. Every added feature that isn't causally connected to the thing you're trying to detect is pure noise the model has to work around, and every added feature that's a near-duplicate of an existing one over-weights that one signal without the model knowing it's being double-counted. Two features that are each genuinely, causally connected to what compromise actually does to a device (it takes longer to respond, and what it sends back changes) beats five features where three are either constant, redundant, or uncorrelated with the thing being modeled.

5. Why This Beats Signatures and Rules for This Threat

Signature- and rule-based detection (Snort/Suricata rules, WAF rulesets, YARA, IOC blocklists) work by matching known patterns: specific byte sequences, specific request structures, specific known-bad indicators. They are fast, precise, and cheap to run, and for a threat that is fully known and fully catalogued, they are genuinely excellent. Their fundamental limitation is that they can only catch what someone has already seen, written a rule for, and shipped to your ruleset. A new CVE, an obfuscated variant of an old payload, or a single byte change to evade a literal string match all slip through a signature engine that was never told to look for that exact shape.

Behavioral anomaly detection doesn't ask what the request looks like on the wire. It asks what the request did to the device: how long did it take to respond, and what came back. That question doesn't require knowing anything about the specific exploit in advance. A genuinely novel technique against the same class of device, one this project's detector has never been told about, still shows up as elevated processing latency if it actually costs the device anything to execute, because the detector isn't looking for the technique, it's looking for the effect.

Signature / Rule-Based

Active Behavioral ML (this approach)

Requires prior knowledge of the specific attack	Requires no prior knowledge of the specific attack
A byte-level variant evades detection	A byte-level variant produces the same behavioral footprint and is still caught
Ruleset needs constant updates as new CVEs appear	Detects new CVEs against the same device class with zero ruleset changes, as long as they produce a similar effect
Fast and precise for exactly what it's told to look for	Slower to say exactly which CVE, but doesn't need to be told in advance
Blind to genuinely novel technique classes	Catches genuinely novel technique classes that share the same underlying behavioral cost

This project's own validation is direct evidence for the argument, not just theory: the platform's CVE set was built independently and happened to land on exactly the five CVEs Check Point Research later documented in a live Iran-attributed campaign. But the detector's ability to catch that campaign's traffic doesn't come from knowing those five CVEs specifically; it comes from those five CVEs all producing the same underlying signal, elevated device processing time, that the model was built to catch regardless of which specific exploit produced it. A sixth, undisclosed CVE against the same vendor class would very likely be caught the same way, without a single line of detection logic changing.

6. Active Probing, Not Passive Monitoring

A meaningful design distinction, easy to miss: this engine is not a passive network tap. Traditional network intrusion detection (Suricata, Zeek, a SPAN-port packet capture) sits inline or on a mirrored port and inspects traffic as it happens to pass by. If the attacker's traffic doesn't cross that specific vantage point, or if it's encrypted in a way the tap can't inspect, a passive system sees nothing.

This engine instead actively probes every device on a fixed cadence, sending the same kind of ordinary request a legitimate client would send, to the same safe endpoints, and measuring the full round-trip: how long the device took to respond, and what it sent back. This is closer in spirit to synthetic uptime monitoring than to a packet-sniffing IDS. It has two direct consequences worth stating plainly:

- No inline deployment, tap, or mirror port is required. The monitor only needs ordinary network reachability to the device, the same as any legitimate client already has, which matters a great deal in the ownership reality described in the companion threat hunting playbook: for a defender who doesn't administer the target device, there is no tap to install, but there is still a request to send and a response to measure.
- The signal comes from the device's own operating condition, not from inspecting the attacker's packets. Even if the attacker's traffic approaches from a network path the monitor has no visibility into whatsoever, the effect of a successful exploit, a device working harder than it should, still shows up directly in how that device responds to the

monitor's own ordinary probe. The detector is reading the device's condition, not trying to catch the attacker mid-transit.

This is the practical meaning of "behavior, responses, and conditions" as the object of detection, rather than packet contents: the engine cares what state the device is in right now, evidenced by how it behaves when asked an ordinary question, not what bytes happened to be visible on some segment of wire.

7. Where Signatures and Rules Still Matter

None of the above is an argument that signature- and rule-based tools should be retired. For a fully known, fully disclosed threat where speed and precision matter more than generalization, a signature match is faster, cheaper, and produces zero ambiguity: it either matched a known-bad pattern or it didn't. IOC blocklists, WAF rules for disclosed CVEs, and YARA signatures for known malware families remain the right tool for exactly that job, and a mature defensive posture runs both layers together rather than choosing one. The case made in this document is narrower and specific: for the open-ended, continuously evolving reconnaissance-and-exploitation activity this project models, where new variants and undisclosed techniques against the same device class are the realistic ongoing threat, behavioral detection catches what signatures structurally cannot, and that is the gap this architecture was built to close.

8. Summary

Isolation Forest was chosen because it's linear-time, needs no labeled attacks, handles multiple features jointly without distance or density computation, and produces an explainable per-reading result. The two-feature design is a deliberate consequence of testing what each candidate features actually contributed, not an oversight. Active probing over passive monitoring means the detector needs no special network position and reads device condition directly rather than depending on visibility into attacker traffic. And the behavioral approach as a whole catches what signature and rule-based tools structurally cannot: variants, obfuscation, and genuinely novel technique classes that were never explicitly catalogued in advance, which is precisely the shape of the threat this platform is built around.