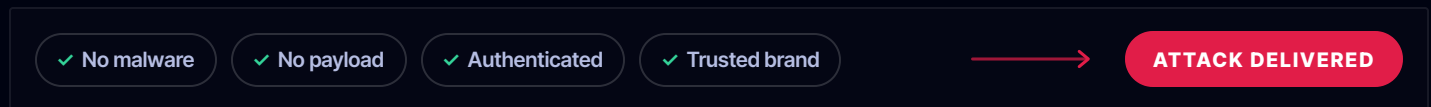


THE PROBLEM

AI is breaking email security **three ways at once.**

Attackers now use AI to **scale, disguise, and personalize** every message. Legacy email security only recognizes attacks it has seen before, and with AI, it never sees the same attack twice.



VOLUME · AUTOMATION

100×

UNIQUE ATTACKS, ONE PROMPT

AI turns a single prompt into **100 unique messages** in seconds, at near-zero cost. No two share enough to fingerprint.

CRAFT · EVASION

4×

EVASION TECHNIQUES STACKED

The average attack now layers **four evasion techniques**, drawn from 36 we track, to slip past each control in turn.

PERSONALIZATION · AIM

3.6%

OF VOLUME, 98% OF THE RISK

The precision-personalized sliver, just **3.6% of volume, carries 98% of the financial loss**. The rare attack is the expensive one.

Volume, craft and personalization: StrongestLayer Phishing Genome study and internal threat research, 2025–26.

54%

And it works. AI-written phishing is now clicked **54% of the time**, matching a skilled human attacker, up from ~12% for traditional mail. The craft is no longer the bottleneck.

Harvard (Heiding & Schneier), Harvard Business Review 2024; 54% vs 12%, Microsoft Digital Defense Report 2025.

Every prior generation, **signatures, reputation and machine learning**, has to see an attack, often several times, before it learns to stop it. AI makes sure no two messages are ever alike, so there is never a second sighting to learn from.

You cannot filter an attack that has never existed. **You have to reason about what it is trying to do.**

THE SOLUTION

An agentic reasoning platform. It reads intent, not patterns.

For every email, autonomous agents reason across four lenses, **intent, harm, anomaly and deception**, and read the attack's genome: what it wants, how it evades, how precisely it is aimed. In context, against how your business actually operates.

REASON ACROSS FOUR LENSES

Intent

Harm

Anomaly

Deception

TO READ THE ATTACK GENOME

Intent · what it wants

Evasion · how it hides

Personalization · how aimed

One engine reads both jobs: the **lure aimed at your people** and the **instruction aimed at the AI** reading your mail. **No competitor reads the second.**

THE OUTCOMES

95%

FEWER FALSE POSITIVES

Real-world FP near 0.5%, where legacy tools run in the tens.

<2 min

ALERT RESOLUTION

From 20+ minutes to under two.

15 min

TO PROTECTION

API-based. No rip and replace, no MX change.

WHO WE WIN AGAINST, AND WHY

Proofpoint · Mimecast · Barracuda

GEN 1 GATEWAYS

Signatures and reputation. The generation already at your MX that these attacks are designed to pass.

Abnormal

BEHAVIORAL AI

A model of "normal" is guessing when the mail is built to look normal, and it never reads the instruction aimed at your AI.

Sublime

RULES-AS-CODE

A rule only catches an attack already seen, and you staff a detection engineer to write it.

StrongestLayer reasons on intent, covers **both the human lure and the AI instruction**, and needs no rules to write. That is the architecture the others do not have.

See what your current stack is missing.

Built by the operators who built the last generation of email security: Proofpoint, FireEye, McAfee, Google / Mandiant.

[Run a threat assessment →](#)strongestlayer.com/contact