

Product Technical Overview

Customer Edition

AutoPurple Security Technology Corp.

Document AP-DOC-PUB-2026-002 · Public · Revision 1.0 · 2026

Contents

Executive Summary	2
1. Company and Approach	2
2. Product Portfolio at a Glance	2
3. TALON — Threat Deception and Attribution System	3
3.1 Positioning	3
3.2 Design principles	3
3.3 What TALON delivers	3
3.4 Compliance boundary	3
3.5 Not disclosed	3
4. SAOAS — Autonomous Security Decision Orchestration Platform	4
4.1 Positioning	4
4.2 Controlled autonomy	4
4.3 Architecture, in concept	4
4.4 Core capabilities	4
4.5 Resilience under graded attack	4
4.6 Compliance positioning	4
5. The TALON and SAOAS Closed Loop	5
6. Argus — External Attack Surface Management	5
6.1 Positioning	5
6.2 Core capabilities	5
6.3 Typical use and engagement	5
7. Spectre — Human and Social-Engineering Assessment	5
7.1 Positioning	5
7.2 Core capabilities	5
7.3 Typical use and engagement	5
8. Mirage — AI System Security	5
8.1 Positioning	5
8.2 Core capabilities	5
8.3 Typical use and engagement	6
9. Aegis — Enterprise Security Operations Center	6
9.1 Positioning	6

9.2 Core capabilities	6
9.3 Typical use and engagement	6
10. Capabilities — The Full Purple-Team Practice	6
11. Deployment, Compliance, and Data Handling	6
12. Engagement Model	7
13. Intellectual Property and Assurance	7
14. Contact	7

Executive Summary

AutoPurple Security Technology Corp. develops products for active defense and security automation on a foundation that is legal, controlled, and auditable. Our work is organized around two flagship platforms — TALON, for threat deception and attribution, and SAOAS, for autonomous security decision orchestration — supported by a four-product suite that covers the external attack surface, human factors, AI-system security, and enterprise security operations.

Every capability we ship originates from a systematic body of adversarial security research and is engineered into a product that a defender can deploy, govern, and audit. This document is a customer-facing overview of that product portfolio. It describes what each product does, how the products relate to one another, how they are deployed, and how we work with clients. It does not disclose implementation methods, detection or fingerprinting algorithms, decision-model parameters, positioning mathematics, or other technical secrets; those are protected as trade secrets and are addressed only under a signed agreement during an engagement.

1. Company and Approach

AutoPurple is a purple-team security company. The name reflects our founding idea: the offensive team's insight (red) and the defensive team's discipline (blue) belong together, unified as controlled autonomy. We build systems that see a network the way an attacker does, and then deliver that understanding to the defender as operational certainty — without ever leaving the defender's own environment.

Three commitments run through every product:

- **Legal, controlled, auditable.** Automation and active defense operate strictly within the client's authorized environment, under tiered authorization, with a complete and tamper-evident audit trail.
- **Control before autonomy.** Our systems are designed so that routine actions can proceed automatically within preset boundaries, sensitive actions require confirmation, and out-of-bounds actions are prohibited by design.
- **Evidence over claims.** Capability is stated to the standard of actual deployment and acceptance. We do not assert more than the evidence supports.

Our products are the engineered result of research programs in AI-system security, adversarial machine learning, attacker attribution, and human-factor security. That research lineage is the reason each product line rests on a defined mechanism rather than on ad-hoc tooling.

2. Product Portfolio at a Glance

The portfolio has two layers that reinforce each other.

Flagship platforms carry the core of the practice:

- **TALON** — Threat Deception and Attribution System. The intelligence source and attribution hub of active defense.

- **SAOAS** — Autonomous Security Decision Orchestration Platform. The controlled decision and response core for security operations.

The four-product suite provides full-stack, AI-native coverage across the domains a modern security program must address:

- **Argus** — External Attack Surface Management.
- **Spectre** — Human and Social-Engineering Assessment.
- **Mirage** — AI System Security.
- **Aegis** — Enterprise Security Operations Center (XDR with an AI-native SOC).

Each product stands on its own and can be adopted independently. Adopted together, they form a single, self-reinforcing program: the attack surface feeds the operations center, the operations center consumes attribution intelligence, and deception informs both.

3. TALON — Threat Deception and Attribution System

3.1 Positioning

TALON is the intelligence source and attribution hub of an active-defense program. It is built for organizations that can already run security operations and that have a genuine need to move beyond “we detected an intrusion” toward “we understand this adversary and where they come from.” TALON constructs a layered deception environment inside the defender’s own network, draws attackers into that controlled space, and profiles them from their own interactions.

3.2 Design principles

- **Legal first, never out of bounds.** All collection and attribution take place inside the defender’s own network. Anything that could touch a network boundary is placed under authorization and audit. TALON does not probe, reach into, or attack external systems.
- **Passive harvesting.** The system does not send probes to external targets. It waits for and induces attacker interaction, and gathers evidence from that interaction. This is both a compliance property and the reason the resulting intelligence carries evidentiary value.
- **Truth first, evidence-driven.** Attacker profiling and origin judgments rest on a traceable chain of evidence. Uncertainty is labeled with confidence and physical limits; positioning precision is never overstated.

3.3 What TALON delivers

- Conversion of isolated attack indicators into a stable operator profile.
- Source attribution that converges by tiers — from region, to city-level network segment, to facility-level — each annotated with a confidence level and its physical bounds.
- High-confidence adversary intelligence suitable for operational decisions and, where appropriate, for coordination with public-safety partners.

3.4 Compliance boundary

TALON’s entire design is organized around achieving effective active defense strictly inside legal limits. Active measurement, where used, is standard network ranging and is distinct from intrusion or out-of-bounds counter-attack. The system is not a tool for “hacking back.”

3.5 Not disclosed

The detection and fingerprinting techniques, the attribution and convergence methods, and the deception construction are protected as trade secrets and are not described here. Positioning precision is expressed as tiered convergence with confidence, not as a numeric guarantee.

4. SAOAS — Autonomous Security Decision Orchestration Platform

4.1 Positioning

SAOAS is the autonomous decision and response orchestration core for security operations. It addresses the structural problem of modern security operations — alert overload, delayed response, and a human bottleneck — by carrying the real-time triage, decision, and orchestration load, so that a security team can focus on high-value judgment and strategy rather than repetitive review.

4.2 Controlled autonomy

SAOAS does not pursue unconstrained automation. Its autonomy is bounded, governable, and revocable. The system acts on its own within an authorization boundary set by the operator, and the operator retains the ability to intervene, pause, and reclaim control at any time. Autonomy deepens in stages as the operator's trust is established, and always at the operator's discretion.

4.3 Architecture, in concept

SAOAS is organized in four cooperating layers:

- **Sensing and real-time triage** — distinguishes attack traffic from legitimate business at the point of entry, at low latency, filtering noise and surfacing real threats.
- **Modular decision engine** — dozens of single-purpose decision modules covering detection, correlation, classification, and disposition, coordinated through one framework and independently extensible.
- **Orchestration and execution** — turns decisions into response actions, binding each execution to the authorized decision that produced it, so that no action proceeds without provenance.
- **Governance and control** — tiered authorization, boundary control, and end-to-end audit that run through every layer.

4.4 Core capabilities

- Real-time discrimination of attack versus legitimate traffic at the point of entry.
- A modular decision engine that is composable, maintainable, and extensible as threats evolve, without disturbing overall stability.
- Adaptive response orchestration whose autonomy depth advances in stages under authorization.
- A three-tier authorization model: routine actions may run automatically within preset limits; sensitive actions require operator confirmation; out-of-bounds actions are prohibited outright, with dual confirmation on high-impact actions.
- Explainable decisions and a complete, tamper-evident audit trail for compliance review and post-incident analysis.

4.5 Resilience under graded attack

SAOAS is built to remain orderly and controlled across attack intensities. Under alert storms it keeps triage stable and prevents high-value events from being buried. Against targeted evasion it cross-checks across multiple modules rather than relying on a single dimension, and where confidence is insufficient it defers to human confirmation rather than acting on uncertainty. Against highly capable adversaries it degrades to human control by design — the operating principle is “return to human control rather than execute out of control.” Resilience here does not mean automatic immunity to every attack; it means that triage stays orderly, action stays controlled, and a human can intervene at any moment.

4.6 Compliance positioning

SAOAS is a decision and orchestration tool for security operations. Its autonomy serves defense within the operator's own environment and is subject to tiered authorization and audit. Its controllability design lets it meet the requirements that serious environments place on an automated security system for compliance and auditability.

5. The TALON and SAOAS Closed Loop

TALON and SAOAS are complementary and, together, form a closed loop. TALON profiles the adversary and traces the source inside the defender's own deception environment, producing a high-confidence adversary profile. SAOAS takes that profile as input for controlled autonomous decision and response, and can in turn task TALON's deception resources to gather fuller intelligence. The two together join "seeing the adversary clearly" and "acting efficiently" into a single automated purple-team loop: deception and attribution, decision orchestration, and response.

6. Argus — External Attack Surface Management

6.1 Positioning

Argus continuously maps what an organization exposes to the outside world, so that the defender sees the attack surface before an attacker does.

6.2 Core capabilities

- Continuous discovery and mapping of exposed assets, ports, and services.
- Certificate and credential-exposure monitoring.
- Attack-surface snapshots over time, with prioritized risk findings.

6.3 Typical use and engagement

Argus suits organizations that need an authoritative, always-current view of external exposure across a growing estate. It is licensed modularly by asset scale, and priced after an independent evaluation.

7. Spectre — Human and Social-Engineering Assessment

7.1 Positioning

Spectre addresses the human layer of security, where most real incidents begin. It combines assessment, awareness training, and a measurable feedback loop.

7.2 Core capabilities

- AI-generated social-engineering assessments tailored to the organization, rather than static templates.
- Coverage of physical social-engineering and AI-deepfake scenarios, not only email phishing.
- Awareness curriculum and an NPS-style feedback loop that make outcomes measurable.

7.3 Typical use and engagement

Spectre suits organizations that want realistic, current social-engineering assessment and durable behavior change across the workforce. It is licensed modularly by headcount, and priced after an independent evaluation.

8. Mirage — AI System Security

8.1 Positioning

Mirage secures AI-native business. As organizations put large models and autonomous agents into production, Mirage provides the red-teaming, monitoring, and compliance layer those systems require.

8.2 Core capabilities

- Large-language-model red-teaming.
- Behavior monitoring for autonomous agents.
- Compliance assessment aligned to frameworks such as the EU AI Act and others.

8.3 Typical use and engagement

Mirage suits organizations deploying AI and agent systems that must be secured and governed under a real regulatory context. It is licensed modularly by the number of AI systems, and priced after an independent evaluation.

9. Aegis — Enterprise Security Operations Center

9.1 Positioning

Aegis is the enterprise security operations command center: an XDR platform with an AI-native SOC that governs the rest of the suite. It brings endpoint, network, and intelligence together with automated response into one operational picture.

9.2 Core capabilities

- A full capability stack spanning endpoint detection, network analysis, threat intelligence, and automated response.
- Consolidation of the suite's signals into a single command center.
- An AI-native detection and response core.

9.3 Typical use and engagement

Aegis suits organizations that want a unified command center rather than a patchwork of point tools. It is delivered as a private deployment, scoped to the client's environment, and priced accordingly.

10. Capabilities — The Full Purple-Team Practice

Delivered across the product suite and through bespoke engagements, AutoPurple's practice covers the disciplines a modern security program relies on:

- Threat Hunting
- Incident Response
- Digital Forensics
- Threat Intelligence
- Attacker Attribution
- Threat Deception
- Red and Purple Teaming
- Penetration Testing
- Social-Engineering Assessment
- LLM Red-Teaming
- Attack Surface Management
- Security Automation (SOAR)
- Detection Engineering
- Vulnerability Assessment
- Compliance and AI Governance

11. Deployment, Compliance, and Data Handling

AutoPurple products are delivered primarily as private deployments, with a modular, pluggable architecture that is configured to the client's environment and requirements. We provide deployment, customization, and ongoing operations support.

All deception and attribution derive from adversary interaction inside the defender's own environment; nothing reaches beyond the defender's network boundary. Active measurement, where used, is standard network ranging and is distinct from intrusion. Every sensitive action is placed under tiered authorization, and all key decisions and executions are recorded in a tamper-evident audit trail. Client submissions and engagement data

are treated as confidential and handled under a signed agreement.

12. Engagement Model

AutoPurple operates a high-touch, one-company-one-engineer model. We keep no standard price list; each engagement is scoped and priced from the client's real needs and security posture. The path is consistent:

1. **Independent red-team analysis.** A consulting engineer performs an independent security analysis of the client's environment to establish the real attack surface.
2. **Scoping and negotiation.** From the findings and the client's business reality, the required modules and service boundaries are agreed.
3. **Posture verification.** The real security posture and adversary intensity are verified as the objective basis for the plan.
4. **Bespoke pricing.** Pricing is set to the actual modular scope — one company, one price — and settled by formal contract.

The modular architecture is precisely why every engagement differs, and why an independent evaluation comes before any plan or pricing.

13. Intellectual Property and Assurance

AutoPurple protects its core technical methods as trade secrets, rather than trading them for the time-limited disclosure of a patent, and registers its software under copyright. We hold to an evidence-driven posture: product capability is described to the standard of actual deployment and acceptance, and never beyond the evidence.

14. Contact

To begin, request an evaluation. An independent assessment by our engineer and officer precedes any plan or pricing.

Evaluation: martincto@autopurple.org

Web: <https://autopurple.org>

AutoPurple Security Technology Corp. — active defense and security automation.

This document is a product-concept and architectural overview for prospective clients. It does not contain implementation methods, decision-model parameters, detection or fingerprinting algorithms, positioning mathematics, or other trade secrets. Stated capability is to the standard of actual deployment and acceptance within legally authorized scope. Product technical whitepapers are controlled-distribution materials provided separately following engagement.