

AutoPurple 产品与市场说明

商业版

奥紫安全技术 (AutoPurple Security Technology Corp.)

文件编号 AP-DOC-PUB-2026-003 · 公开 · 第 1 版 · 2026

引言

AutoPurple (奥紫) 是一家网络安全公司,做的是“主动防御”和“安全自动化”两件事。用最朴素的话说:过去的安全产品像被动的门锁和报警器,坏人来了才响;我们做的东西,是让防守方能像攻击者一样看清自己的网络,在被攻破之前先发现问题,并且在真被攻击时,不只知道“有人闯进来了”,而是弄清楚“这个攻击者是谁、从哪来”,同时让机器在授权范围内自动、可控地把大量重复的安全工作扛下来,把人从告警的海洋里解放出来。公司名 AutoPurple (自动化紫队)就来自这个想法:把攻击方的视角(红)和防守方的纪律(蓝)合在一起,变成“可控的自动化”。

有一条底线贯穿我们所有产品:一切能力都只在客户自己的授权环境内运行,合法、可控、可审计,绝不越出客户网络的边界。

1. 产品效果 (客户用了之后,得到什么)

我们的产品线由两个旗舰平台加四个专项产品组成,合起来解决企业安全里最要命的几个问题。用大白话讲,客户能得到的效果是这样几条。

看清自己在外暴露了什么。很多被攻破的企业,其实自己都不清楚对外开了哪些门、哪些服务、哪张证书快过期。我们的产品持续地、自动地帮客户把这些“暴露面”摸清楚,在攻击者发现之前先让防守方看到,并按风险高低排好先后。

知道攻击者是谁、从哪来。大多数安全产品只能告诉你“被攻击了”。我们的旗舰平台能在客户自己的网络里布下一层“诱捕环境”,把攻击者引进这个受控空间,从他自己的行为里把画像刻出来,把零散的攻击线索一步步收敛成“区域到城市网段再到具体设施”这样越来越精确的来源判断,而且每一步都标明可信程度和物理上的边界,不夸大。

把安全团队从告警海里解放出来。现代安全团队最大的痛点是告警太多、响应太慢、全靠人扛。我们的另一个旗舰平台,让机器在授权范围内自动完成实时分诊、判断和处置这些重活,人只做最有价值的判断和决策。机器扛量,人扛关键。

管住“人”和“AI”这两个新的薄弱点。绝大多数真实的安全事故,其实是从“人”这一环被突破的(被骗、被钓鱼)。我们有专门的产品做贴合企业实际的社工演练和安全意识训练,并且能把效果量化出来。而随着企业越来越多地用上大模型和 AI 助手,这些 AI 系统本身也成了新的攻击面,我们有专门的产品为它做红队测试、行为监控和合规评估。

一个统一的安全指挥中心。我们有一款企业级的安全运营中心产品,把终端、网络、情报和自动响应汇到一张作战图上,让客户不再是东一个工具、西一个工具地打补丁,而是有一个统一的、带 AI 的指挥中枢。

这些产品每一个都能单独使用;合在一起,则形成一个互相加强的闭环:暴露面的发现喂给运营中心,运营中

心消费溯源情报，诱捕又同时给两边提供养料。客户用得越全，整体效果越接近“一加一大于二”。

2. 产品构成（产品由什么组成）

整个产品体系分两层，互相支撑。

上层是两个旗舰平台，是整套实践的核心：

- **威胁诱捕与溯源系统 (TALON)**：主动防御的情报来源和溯源中枢。它在防守方自己的网络里构建分层的诱捕环境，把攻击者引入受控空间，从其交互中刻画攻击者画像、追溯来源。
- **安全决策自主编排平台 (SAOAS)**：安全运营的受控决策与响应核心。它承担实时分诊、判断和编排这些负荷，让团队专注于高价值判断。

下层是四个专项产品，覆盖一个现代安全体系必须照顾到的四个领域，并且都是“AI 原生”的（从底层就以 AI 为核心构建，而不是在老架构上贴一层 AI 外壳）：

- **外部攻击面管理 (Argus)**：持续测绘企业对外暴露的资产、端口和服务。
- **人因与社会工程评估 (Spectre)**：针对“人”这一层的评估、意识训练和可量化的反馈闭环。
- **AI 系统安全 (Mirage)**：为企业上线的大模型和智能体提供红队测试、行为监控和合规评估。
- **企业安全运营中心 (Aegis)**：带 AI 原生 SOC 的 XDR 平台，统管全套产品，把各路信号汇成一张作战图。

以其中的安全决策平台 (SAOAS) 为例，它在设计上分四个协同的层：入口的实时感知与分诊，在流量进来的第一时间就把攻击和正常业务分开；模块化的决策引擎，几十个各司其职的决策模块由一个框架统一协调、可独立扩展；编排与执行，把决策变成响应动作，每一次执行都绑定到产生它的那条授权决策上，没有来历的动作不会执行；以及贯穿全程的治理与控制，也就是分级授权、边界控制和端到端审计。

交付上，我们以私有化部署为主，采用模块化、可插拔的架构，按客户的实际环境和需求来配置，并提供部署、定制和长期运维支持。

3. 产品可靠性（为什么可以放心用，尤其关系到用户体验）

对我们这类产品，“可靠”不是一句口号，而是直接关系到客户的正常业务会不会被误伤、系统会不会失控。我们从设计上就把可靠性放在自主能力之前，具体体现在下面几点。

可控优先于自治。系统的自主是有边界、可治理、可随时收回的。例行的、在预设范围内的动作可以自动执行；敏感的动作必须由操作者确认；越界的动作从设计上就被禁止，影响重大的动作还要双重确认。操作者在任何时刻都可以介入、暂停、收回控制权。自主的深度是随着信任一步步放开的，而且始终由客户说了算。

不误伤正常业务。系统在流量进入的第一时间，就以很低的延迟把攻击流量和正常业务区分开，过滤掉噪音，只把真正的威胁浮上来。这一条直接关系到用户体验：它保证的是正常客户不会被挡在门外，而真正的威胁也不会被淹没。

越是被猛攻，越要稳得住。面对告警风暴，系统保持分诊有序，不让高价值事件被埋没；面对针对性的规避手法，它靠多个模块交叉验证，而不是把宝押在单一维度上；当可信程度不够时，它宁可交给人工确认，也不在不确定的情况下贸然行动；面对能力极强的对手，它的设计原则是“宁可退回人控，也不失控执行”。这里说的“可靠”，不是号称对一切攻击自动免疫，而是保证：分诊始终有序、动作始终受控、人可以随时介入。

决策可解释、全程可审计。系统的每一个判断都能说清楚为什么这么判，所有关键决策和执行都记录在防篡改的审计链里，既能满足合规审查，也方便事后复盘。

实事求是，不夸大。我们对能力的描述，只说到实际部署和验收的标准，绝不说超出证据支持的话；溯源的精度用“分级收敛加置信度”来表达，而不是给一个做不到的数字保证。

4. 市场效果（市场机会、定位与商业模式）

市场需求来自几个结构性的、短期内不会消失的痛点：企业对外的暴露面越来越大、越来越难摸清；安全告警的数量早已超过人能处理的极限，人力成了瓶颈；绝大多数真实事故从“人”这一环被突破；而企业大规模用上 AI 和智能体之后，又多出一整块全新的攻击面和合规压力。这几件事叠加在一起，决定了“安全自动化”“AI 安全”“主动防御”是整个网络安全行业里增长最快、也最被需要的方向之一。

在这个方向上，我们的定位有几个别人不容易复制的差异点。

紫队一体。大多数公司要么偏红（攻），要么偏蓝（守），我们把两者合成一个“可控自治”的整体，既能像攻击者一样看问题，又交付给防守方可以直接用的确定性。

合法、可控、可审计。我们所有的主动防御和自动化都严格在客户授权环境内、分级授权、全程留痕，这让我们能进入那些对合规和可审计性要求很高的严肃行业。

一司一专家的高端定制。我们不设标准价目表。每一次合作，都先由我们的工程师对客户环境做一次独立的红队分析，摸清真实的攻击面，再据此和客户商定需要哪些模块、服务边界到哪里，最后按实际的模块范围一司一价、以正式合同结算。

AI 原生。我们的产品不是在老架构上补一个 AI 功能，而是从底层就以 AI 为核心来构建。

商业模式上，我们以私有化部署交付为主，按资产规模、人数、AI 系统数量等维度模块化授权，评估后独立定价。正因为架构是模块化的，每一次合作的范围和价格都不一样，这也是为什么任何方案和报价之前，都要先做一次独立评估。

知识产权上，我们把核心技术方法作为商业秘密来保护，而不是用专利那种有时限的公开去换保护，软件部分登记软件著作权。这样既守住了技术护城河，又能对客户如实交付。

结语

AutoPurple 做的是一件朴素的事：让防守方看得清、反应得快，同时始终合法、可控、可审计。我们对能力的每一句描述，都以真实的部署和验收为准。合作的第一步，是请我们的工程师和负责人为您做一次独立评估；任何方案和报价，都在评估之后。

评估入口：martincto@autopurple.org 网站：https://autopurple.org

奥紫安全技术（AutoPurple Security Technology Corp.）——主动防御与安全自动化。