

基于关联分析的 APT 态势自主觉察方法

孟楠¹ 崔泉飞¹ 石悦¹

MENG Nan CUI Xiaofei SHI Yue

摘要

高级持续性威胁攻击在网络空间博弈战场中占据重要位置,其隐蔽性、持久性、复杂性等特点使得传统检测方案无法取得良好的检测效果。针对此,文章提出一种基于关联分析的 APT 态势自主觉察模型,通过深度分析现有的威胁情报资源成功构建用于 APT 态势感知的综合性情报图谱,通过实时比对多源数据关联产生的攻击图与情报图之间的相似性,实现高敏感、实时性的 APT 态势自主觉察。公开 APT 攻击数据集下平均 98% 的准确率和秒级的预警时间证明了所提方法的有效性和可用性。

关键词

高级持续性威胁;关联分析;态势觉察;情报图谱

doi: 10.3969/j.issn.1672-9528.2025.08.021

0 引言

在网络空间博弈的第五战场中,高级持续性威胁(advanced persistent threat, APT)一直是重要组成部分,具有国家背景支持、针对性强、隐蔽性高、持续性久、组织严密、技术复杂等特点。图 1 为近年来具备一定代表性的 APT 攻击事件列表。根据 360 发布的《2023 年全球高级威胁(APT)研究报告》^[1],2023 年累计公开发布 APT 报告 731 篇,涉及 APT 组织 135 个,其中首次披露的 APT 组织 46 个。从全球范围看,各个国家的政府、国防军工、信息技术、外交、教育等均受到了 APT 攻击的严重影响,而我国作为主要的受害国家,正面临着尤为严峻的网络攻击威胁^[2],对 APT 攻击的检测和防御刻不容缓。

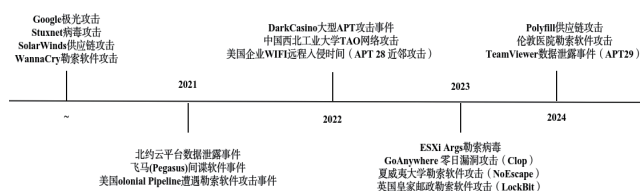


图 1 近年来典型 APT 攻击事件

传统的 APT 攻击检测手段高度依赖于专家知识和人工设定的标准,通过精心设计的检测规则与特征库(如特征值检测技术、校验和检测技术、启发式检测技术以及主动防御技术^[3-15])来实施。然而,上述方法在面对 APT 攻击组织的高超技术和丰富资源时存在显著的局限性:难以处理攻击跨度时间长的数据;难以进行实时检测;无法应对 0 day 漏洞等。尽管 APT 攻击行为以其持久的持续性和高度的隐蔽性著称,但不同 APT 攻击组织在各类安全事件中展现出的技战术手段

及所利用的攻击工具却呈现出一定的共通性。例如,海莲花(APT-C-00)组织频繁采用鱼叉式网络钓鱼攻击与水坑攻击策略,并擅长利用及改造公开的攻击工具与开源项目以增强其攻击效能。鉴于此,本文提出一种基于关联分析的 APT 态势自主觉察模型,旨在通过深度分析与学习现有的威胁情报资源,以精准把握 APT 攻击活动的特征模式与行为规律,进而实现对 APT 攻击活动的实时态势感知与有效监控。具体来说,本文的主要贡献如下:

(1) 创新性地将杀伤链模型、MITRE 的威胁知识库以及 360 情报中心等权威威胁情报中心提供的 APT 情报数据进行深度融合,构建了用于 APT 态势感知的综合性情报图谱。

(2) 通过实时比对由安全设备日志、网络流量等多源数据关联分析生成的攻击图与情报图之间的相似性,实现高敏感、实时性的 APT 态势自主觉察。

(3) 在公开的 APT 攻击数据集上对本文所提方法的可用性和实时响应能力进行了验证。

1 基于关联分析的 APT 态势觉察模型

基于关联分析的 APT 态势自主觉察模型如图 2 所示,该模型的最终目的是探测网络内部是否存在疑似 APT 组织发起的攻击,以便在攻击者成功窃取数据或破坏系统之前,能够及时发出预警并采取相应的防御措施。本模型包含 APT 情报图构建、攻击路径识别、APT 态势预警 3 个核心模块。

(1) APT 情报图构建:从 MITRE 以及 360 情报中心等威胁情报平台的数据源中收集 APT 相关的基础情报数据,采用文本分析的方法从中提取 APT 组织、攻击技术、攻击工具、攻击对象、漏洞、IOC 情报、网络设备等基本实体以及实体间关系,从而形成反映不同 APT 组织的攻击模式和演进路径的知识图谱,作为态势预警的领域知识支撑。

1. 中国信息通信研究院 北京 100083

(2) 攻击路径识别: 基于时序关系从组织内部各类安全设备上报的安全事件日志中形成网络告警时序链, 基于流量数据提取组织内部的网络通信关系图, 根据各个告警相关的网络设备间是否通联对告警时序链进行初步拆分, 根据基于专家经验构建的攻击模式知识库识别攻击路径。

(3) APT 态势预警: 鉴于攻击路径图中的节点和边与 APT 情报图中的节点和边可能存在不对等的情况, 模块首先需要将攻击路径图中的节点和边转换为情报图能够理解和识别的内容, 从而建立起两个图之间的精确对应关系。随后, 根据属性、标签、结构的相似性来判断攻击图是否属于情报图, 从而实现 APT 态势的实时预警。

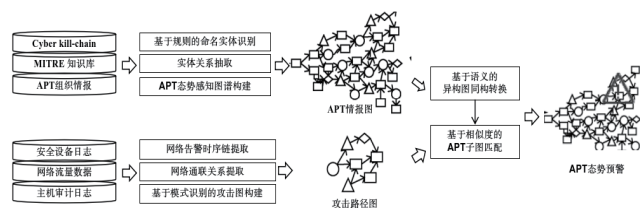


图2 基于关联分析的 APT 态势自主觉察模型

2 基于文本分析和图模型的 APT 情报图构建

表1是本文构建 APT 情报图的基础数据, 通过官网下载、网络爬取等多种手段获得, 都是非结构化的文本。

表1 APT 情报图构建的基础数据

MITRE		威胁情报中心	
类别	涵盖内容	类别	涵盖内容
ATT&CK	攻击技术, 技术描述, 相关工具, APT 组织利用情况	APT 组织	组织名称、ID、来源、最近活动时间、目标国家/地区
CWE	CWE ID, 漏洞描述, 成因, 关联 CVE ID	技战法	常用攻击技术、攻击工具
D3FEND	攻击技术, 攻击对象 (文件、进程等)	APT 情报	IP、HASH、URL、域名等 IOC 情报
CAPEC	攻击模式, 模式描述, 关联漏洞的 CWE ID, 关联 ATT&CK 技术	相关案例	攻击发生时间、攻击目标、攻击工具、攻击手段

虽然表1中数据来自不同的机构, 但属于同一机构的文本在结构上是相似的, 以 MITRE ATT&CK 官网文本为例, 攻击技术、攻击工具、APT 组织等名称都是超链接文本的形式。因此本研究针对不同来源的基础数据设定各自的解析规则 (正则表达式、关键字符切分等) 来识别各种实体 (名称和属性标签)。在实体抽取的基础上, 需要构建实体间的关联关系, MITRE 的知识库中本身就包含了一些映射关系, 如 CAPEC 知识库中包含 [攻击模式 → CWE ID → ATT&CK 技术 ID] 的映射, D3FEND 中又涵盖 [ATT&CK 技术 ID →

攻击对象] 的映射, 融合这些映射关系, 就可以构建攻击技术、攻击对象、漏洞、攻击工具等实体之间的关联; 针对 APT 情报数据中计战法描述的自由文本部分, 本文则根据分词等技术自动识别并标注文本内嵌的各类实体及其相互关联。以图的形式存储上述抽取的实体以及实体间关系, 如图3所示, 是本文最终构建的 APT 态势感知的情报图谱结构。

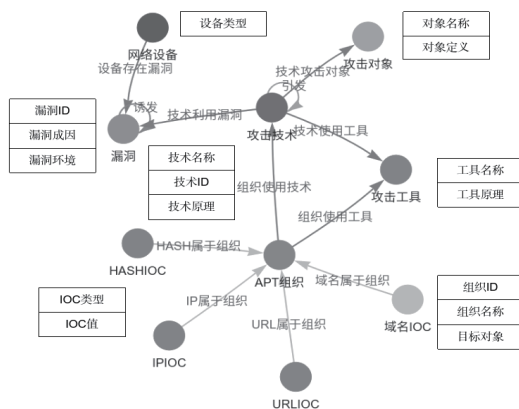


图3 面向 APT 态势觉察的情报图谱结构

3 基于多源数据关联的攻击路径识别

在攻击模式库的监督下, 通过整合安全事件日志和流量数据, 实现了对攻击路径的有效识别, 旨在从海量数据中快速发现值得注意的、重要的威胁告警, 为后续进行 APT 态势感知提供判定输入。

基于多源数据关联的攻击路径识别算法如算法1所述, 包括网络告警时序链构建、网络通信关系图提取、告警时序链逻辑拆分、攻击路径识别4个步骤:

(1) 收集组织内部各类安全设备上报的安全事件日志, 并根据事件的时间戳和类型, 对日志进行排序和分类, 形成网络告警时序链。

(2) 收集组织内部的网络流量数据, 并提取其中通信对、通信协议等属性信息, 构建组织内部的网络通信关系图。

(3) 根据网络通信关系图, 检查告警时序链中各个告警相关的网络设备间是否存在通信关系, 若告警间无通信关系, 则将其拆分为独立的告警子链。

(4) 基于专家经验构建攻击模式知识库, 拆分后的告警子链与攻击模式知识库, 识别潜在攻击路径。

算法1: 基于多源数据关联的攻击路径识别

输入: 安全事件日志集合 logs, 流量数据集 flowData, 攻击模式知识库 attackPatterns

输出: 攻击路径集合 attackPaths

1. 初始化告警时序链集合 alertChains 为空
2. 初始化网络通信关系图 networkGraph 为空
3. 初始化攻击路径集合 attackPaths 为空
// 构建网络告警时序链
4. 对于 logs 中的每一条日志 log:

a. 提取 log 的时间戳 timestamp 和事件类型 eventType

b. 将 log 添加到 alertChains 中相应的时序链中（根据 timestamp 排序）

// 提取网络通信关系图

5. 对于 flowData 中的每一条流量记录 flow:

a. 提取 flow 的源地址 sourceIP 和目的地址 destinationIP

b. 在 networkGraph 中添加从 sourceIP 到 destinationIP 的边

// 告警时序链初步拆分

6. 对于 alertChains 中的每一条时序链 alertChain:

a. 初始化拆分后的告警子链集合 subChains 为空

b. 对于 alertChain 中的每一个告警 alert:

i. 提取 alert 的相关设备 IP 地址 alertIP

ii. 对于 subChains 中的每一个子链 subChain:

A. 如果 subChain 中的最后一个告警的相关设备 IP 与 alertIP 无直接通信关系（在 networkGraph 中无路径）:

①将 subChain 从 subChains 中移除

②将 alert 作为新子链添加到 subChains 中

B. 否则，将 alert 添加到 subChain 中

iii. 如果 subChain 为空，则直接将 alert 作为新子链添加到 subChains 中

c. 更新 alertChains 为拆分后的 subChains

// 攻击路径识别

7. 对于 subChains 中的每一条告警子链 subChain:

a. 初始化匹配成功的攻击模式 pattern 为空

b. 对于 attackPatterns 中的每一个攻击模式 candidatePattern:

i. 如果 subChain 与 candidatePattern 匹配:

①将 candidatePattern 赋值给 pattern

②将 subChain 添加到 attackPaths 中（作为一条攻击路径）

③跳出循环（不再继续匹配其他候选模式）

ii. 如果未找到匹配模式，则记录为未知攻击路径

8. 返回 attackPaths

4 基于图相似性的 APT 态势预警

APT 态势预警模块能够智能地分析攻击路径图，并将其与 APT 情报图进行精确匹配，从而实现对 APT 攻击的实时预警。在该模块中需要解决的一个主要挑战是攻击路径图与 APT 情报图之间的不对等情况，如组织内部上报的安全事件可能是“频繁登录应用系统失败”“钓鱼邮件”，但在图谱中对应的攻击技术节点是“密码猜测”“鱼叉式网络钓鱼”。为解决这一问题，本模块设计了一套语义相关的映射机制，根据节点名称、节点属性的文本相似度将攻击路径图中的节点和边转换为情报图能够理解和识别的标准格式。

如图 4 所示，是基于图相似性进行 APT 态势预警的基本流程。完成攻击图向情报图的同构转换后，根据攻击图中的节点信息在情报图进行检索发现相关的情报子图，随后根据情报子图和攻击图上边的出入、节点的名称、属性、节点之间的结构来比对两个图之间的相似性从而判定发现攻击图是否命中已有 APT 组织的攻击技战术，从而实现实时的 APT 态势预警。

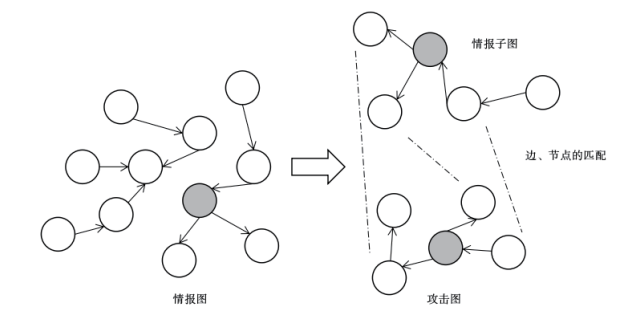


图 4 基于图相似性的 APT 态势预警模型

5 实验与评估

5.1 数据集

本文选择 DAPT2020 数据集进行实验验证，DAPT2020 是由亚利桑那州立大学和美国海军研究实验室共同设计提出的，从自主设计的网络环境中模拟日常使用和高级持续性威胁攻击，然后抓取网络流数据包^[16]。如表 2 所示，是该数据集中高级持续性威胁的攻击阶段和所用攻击技术。

表 2 DAPT2020 数据集信息

攻击阶段	持续时间 /h	使用技术
良性	10	—
侦察	10	网络扫描、应用扫描、账号破解
建立立足点	10	SQL 注入、目录破解、跨站请求伪造、恶意软件下载、命令注入、反向 Shell
横向移动	10	内部扫描、密码抓取、凭证窃取、账号发现、用户账号创建、特权提升
数据泄露	10	数据渗出

5.2 实验设置

APT 态势预警模型本质上是二分类问题，判定对应的网络安全事件是否属于 APT 攻击范畴，因此采用精准率（precision, P ）、召回率（recall, R ）作为评价指标。为验证本文所提模型的有效性和实时性，本文设计了如下两类实验：

- （1）模型有效性实验：随机从 4 种攻击阶段的原始数据抽取 10%、20%、30%、40%、50%、60%、70%、80%、90%、100% 的数据依次加入良性数据中组成 10 组测试数据集，测试模型的识别效果。
- （2）模型实时性实验：修改原始数据的时间戳信息，在良性数据中随机穿插不同攻击阶段的攻击技术数据，测试

模型进行 APT 态势预警的时间和攻击事件发生时间之间的延迟,验证模型是否具备实时响应能力。

5.3 实验结果

模型有效性和可用性实验的结果如图 5 所示,从不同测试数据集下的精准率和召回率来看,平均 98% 的结果表明本模型能够实现精准的 APT 态势预警。同时根据模型预警时间和攻击事件发生之间的时间延迟结果(单位: s)来看,在攻击事件发生后平均 1 s 内模型就能够上报预警,可以看出,一旦出现属于 APT 攻击模式的安全事件,模型能够实时预警,模型具备高度的可用性。

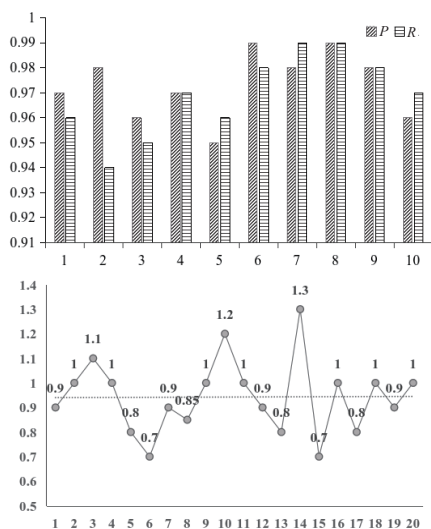


图 5 模型有效性和可用性实验结果

6 总结与展望

本文根据 APT 组织在攻击手段、攻击工具、漏洞利用等战法上的相似性,融合 MITRE、权威威胁情报中心公开的知识数据构建了涵盖 APT 组织、攻击技术、攻击工具、攻击对象、漏洞、IOC 情报、网络设备等多种实体的情报图,用于给 APT 态势觉察提供领域知识依赖。为能在 APT 攻击完全达成其目标之前,有效识别并捕捉其中间阶段的迹象,本文通过整合安全事件日志和流量数据,实现了对攻击路径的有效识别,并将其与 APT 情报图进行精确匹配,从而实现对 APT 攻击的实时预警。通过在 DAPT2020 数据集上的实验验证,本研究证明了所提模型在预警 APT 攻击方面的有效性和实用性,但值得注意的是,模型的性能仍需在更广泛数据集上进行反复验证与迭代优化。

参考文献:

- [1] 360 数字安全.《2023 年全球高级威胁 (APT) 研究报告》[R/OL].2024[2025-01-23].<https://archive.threatbook.cn/threatbook/2023-ThreatBook-CTI-APT-Analysis-Report.pdf>.
- [2] 杨秀璋,彭国军,刘思德,等.面向 APT 攻击的溯源和推理研究综述[J].软件学报,2025,36(1):203-252.

- [3] YE Y F, LI T, ADJEROH D, et al. A survey on malware detection using data mining techniques[J]. ACM computing surveys (CSUR), 2017, 50(3): 1-40.
- [4] WANG R, FENG D G, YANG Y, et al. Semantics-based malware behavior signature extraction and detection method[J]. Journal of software, 2012, 23(2): 378-393.
- [5] SONG W N, PENG G J, FU J M, et al. Research on malicious code evolution and traceability technology[J]. Journal of software, 2019, 30(8): 2229-2267.
- [6] 季一木,张嘉铭,杨倩,等.高级持续性威胁检测与分析方法研究进展[J].南京邮电大学学报(自然科学版),2025,45(1): 1-11.
- [7] 杨秀璋,彭国军,刘思德,等.面向 APT 攻击的溯源和推理研究综述[J].软件学报,2025,36(1):203-252.
- [8] 王邳伟,何晞杰,易鑫,等.基于 APT 活动全生命周期的攻击与检测综述[J].通信学报,2024,45(9): 206-228.
- [9] 张涛,丁伟策,尹冬梅.基于 ATT&CK 模型的 APT 攻击检测防御技术[J].电脑编程技巧与维护,2024(7): 167-169.
- [10] 赵新强,范博,张东举.基于威胁发现的 APT 攻击防御体系研究[J].信息安全学报,2024,24(7): 1122-1128.
- [11] 姚星昆,郑先伟.全流量回溯结合 APT 建模分析技术[J].福州大学学报(自然科学版),2023,51(5): 639-646.
- [12] 谭振江,邹娜,郑月锋.基于 SIEM 系统的 APT 攻击检测框架[J].吉林师范大学学报(自然科学版),2023,44(3): 118-123.
- [13] 陈泽红.基于自适应模糊聚类的无监督 APT 攻击检测方法研究[J].网络安全技术与应用,2023(7): 45-47.
- [14] ALY A, IQBAL S, YOUSSEF A, et al. MEGR-APT: a memory-efficient APT hunting aystem based on attack representation learning[J]. IEEE transactions on information forensics and security, 2024, 19: 5257 - 5271.
- [15] CHEN D W, ZHU P F, YAN H S, et al. APT-DFLC: a defense system framework against APT attack for high security level network based on life cycle[C]//2024 4th Asia-Pacific Conference on Communications Technology and Computer Science (ACCTCS). Piscataway: IEEE, 2024: 730-736.
- [16] MYNENI S, CHOWDHARY A, SABUR A, et al. DAPT 2020-constructing a benchmark dataset for advanced persistent threats[J].Deployable machine learning for security defense, 2020: 138-163.

【作者简介】

孟楠(1982—),女,河南开封人,博士,高级工程师,研究方向:网络和数据安全、ICT 新技术安全领域科研和技术创新、政策和标准制定等。

(收稿日期:2024-12-09 修回日期:2025-07-30)