

SECURITY INSIDER

网安26号院

奇安信网络安全通讯

前沿

奇安信的AI

驱动威胁情报运营实践 P11



P04

月度网络安全态势

P34

从零建设，看通信服务商如何打造
全方位的网络安全“免疫力”

第43期

2024 年 7 月

打造新一代中国特色的安全托管服务

时刻守护您的网络安全



奇安信安全托管服务以威胁发现为基础、以分析处置为核心、以发现隐患为关键、以解决客户问题为目标，7*24小时全天候全方位守护客户网络安全。

两种模式 模式随心选择

- 直营服务模式：奇安信产品+MSS
- 合作服务模式：技术、产品、服务整体托管

多种形态 全面助力建成

- 城市运营中心
- 行业安全运营中心
- 企业安全运营中心

两化融合 帮您真正实现

- 集中服务化：统一监测、预警与通报
- 服务集中化：标准统一、质量可控、资源共享



首创“云地结合”模式

打破传统的托管仅限“云端监测”的认知，开辟“云地结合”模式，云端+地端实现真正的闭环服务。



7*24h实时持续监测

“地球不爆炸，我们不放假”——7*24h持续监测，充分保障常态化运营。



安全事件响应快一步

对APT攻击、网络攻击等长效监测，提前预警，响应快一步。



安全事件处置规范化

事前检测、事中分析、事后总结，真正实现闭环，并在实战中不断加强。



专家“一对一”指导

专家通过视频一对一的方式与您“面对面”沟通指导，不仅让您了解现状，还能防患于未然。

一次 IT 更新引发的全球灾难

7月19日，全球范围内的 IT 系统故障在各国蔓延，导致飞机停飞、医疗系统崩溃、政府和财富 500 强公司的 IT 系统陷入瘫痪。截至格林威治标准时间 18:00，全球超过 3300 个航班被取消，给航空旅行行业造成严重影响。各地的支付系统、银行和医疗保健机构受到了波及。

不过，令人感到震惊的是，引发这次全球 IT 系统重大事故的，不是黑客组织网络攻击，也不是勒索软件，而仅仅是因为一个安全软件更新。网络安全巨头 CrowdStrike 承认，问题是由 Windows 主机上的 CrowdStrike Falcon Sensor 软件自动更新引发的。在 CrowdStrike 发布软件更新后，导致微软 Windows 崩溃并持续蓝屏重启，显示错误消息（俗称“死机蓝屏”）。

CrowdStrike 首席执行官乔治·库尔茨在 X 上的一篇帖子中表示，此次 IT 故障并非由于“安全事件或网络攻击”造成。这意味着黑客、僵尸网络和勒索软件运营商的让全球瘫痪的梦想，由安全巨头 CrowdStrike 实现了。

尽管该公司已撤销这一更新，却对已经发生故障的设备没有作用。安全研究人员警告称，要解决问题需要对每台计算机进行手动修复，这对于拥有数千、数万台计算机设备的企业将会是一项繁重工作。根据 CrowdStrike 公司的最新财报，其客户总数接近 24,000 家，意味着这一事件可能造成广泛影响，成为历史上最大的网络事件之一。

由于问题根源是 CrowdStrike Falcon Sensor 的自动更新。网络安全专家指出，事件凸显了具有远程管理功能的安全工具的风险。在没有事先进行安全检查的情况下安装更新，存在巨大风险。

这一事件可能还会引发安全软件作用的讨论。当前，整个生态系统中以管理权限运行的安全软件太多，现在是时候考虑规避风险的方案了。

此外，软件供应链（包括网络安全供应链）攻击和中断问题同样值得关注，这意味着对软件系统（包括安全软件）推进零信任的必要性。可能即便供应商是“同类最佳”，并不意味着他们就能免疫。

一次软件更新引发全球 IT 事故，这凸显出将业务连续性放在首位的重要性，政企机构需要做好业务弹性规划，以随时应对勒索软件攻击、员工失误或意外 IT 故障的威胁。

总编辑

李建平

2024 年 7 月 1 日



安全态势

- P4 | 国家标准《数据安全技术 个人信息保护合规审计要求》公开征求意见
- P4 | 两部门印发《关于开展“网络去 NAT”专项工作 进一步深化 IPv6 部署应用的通知》
- P4 | 四部门联合印发《国家人工智能产业综合标准化体系建设指南（2024 版）》
- P5 | 国家标准《网络安全技术 关键信息基础设施安全保护能力指标体系》公开征求意见
- P5 | 美议员提出《联邦网络安全法规简化法案》，以改进监管碎片化问题
- P5 | 美国空军部发布《空军部零信任战略》
- P6 | 南非矿业巨头遭网络攻击：被迫隔离 IT 系统 企业运营受干扰

- P6 | 美国大型银行因勒索攻击关闭系统多天，客户无法访问账户或交易
- P6 | 菲律宾国家医保系统泄露超 4200 万用户数据，负责国企高管遭议会公开质询
- P7 | 猴痘病毒爆发期，南非国家卫生实验室因勒索攻击中断服务
- P7 | 供应商被黑导致美国上万家汽车经销商业务停顿
- P7 | 全球叉车巨头科朗遭网络攻击，工厂生产被迫中断超一周
- P7 | AMD、苹果疑遭网络攻击，产品源代码泄漏
- P8 | ServiceNow 多个高危漏洞安全风险通告
- P8 | GitLab 身份认证绕过漏洞安全风险通告
- P8 | Splunk Enterprise 任意文件读取漏洞安全风险通告
- P8 | Rejetto HTTP File Server 模板注入漏洞安全风险通告
- P9 | Artifex Ghostscript 格式字符串漏洞安全风险通告
- P9 | GeoServer 远程代码执行漏洞安全风险通告
- P9 | OpenSSH 远程代码执行漏洞安全风险通告
- P9 | Progress MOVEit Transfer 身份认证绕过漏洞安全风险通告

月度专题



P11

前沿

奇安信的AI

驱动威胁情报运营实践

作为网络安全领域的重要变革力量，人工智能正在重塑组织的威胁情报运营方式。作为国内网络安全的领军企业，奇安信走在 AI 驱动威胁情报运营的创新前沿，正在利用人工智能和大模型的力量为客户提供所需的威胁情报，以更好地保护资产并领先威胁者一步。

安全之道

P34

从零建设，看通信服务商如何打造全方位的网络安全“免疫力”

报告速递

P37

Gartner：软件供应链安全指南

专栏

P56

架构演进叠加客户需求，推动边界安全加速融合

P59

人工智能对密码学的威胁探讨

P63

安全团队指南：
如何创建网络安全的“谷歌地图”

《网安 26 号院》编辑部

主办 奇安信集团

总 编 辑：李建平

安全态势主编：王 彪

月度专题主编：李建平

安全之道主编：张少波

奇安信资讯主编：陈 冲

报告速递主编：刘川琦

专 栏主编：任润波



奇安信集团



虎符智库



安全内参

索阅、投稿、建议和意见反馈，请联系奇安信集团公关部

索阅邮箱：26hao@qianxin.com

地 址：北京市西城区西直门外南路 26 院 1 号

邮 编：100044

联系电话：（010）13701388557

出版物准印证号：内资准印证 京内资准 2123-L0058 号

编印单位：奇安信科技集团股份有限公司

发送对象：奇安信集团内部人员

印刷数量：4500 本

印刷单位：北京博海升彩色印刷有限公司

印刷日期：2024 年 7 月 26 日

版权所有 ©2023 奇安信集团，保留一切权利。

未经奇安信集团书面同意，任何单位和个人不得擅自摘抄、复制本资料内容的部分或全部，并不得以任何形式传播。

无担保声明

本资料内容仅供参考，均“如是”提供，除非适用法要求，奇安信集团对本资料所有内容不提供任何明示或暗示的保证，包括但不限于适销性或适用于某一特定目的的保证。在法律允许的范围内，奇安信在任何情况下都不对因使用本资料任何内容而产生的任何特殊的、附带的、间接的、继发性的损害进行赔偿，也不对任何利润、数据、商誉或预期节约的损失进行赔偿。

内部资料 免费交流

奇安信资讯

P39 | 齐向东出席 2024 中国互联网大会：以 AI 驱动安全 全力护航互联网新质变

P40 | 陕西洞鉴云侦司法鉴定资质获批准

P41 | 北京市人大专题调研组到奇安信安全中心调研

P42 | 齐向东受聘北京交通大学兼职教授仪式暨名师大讲堂活动圆满举办

P43 | 奇安信集团 77 个项目获批教育部第三期供需对接就业育人项目

P44 | 齐向东：发展为先、安全护航 推动数据要素发挥乘数效应

P45 | 豫信电子科技集团董事长李亚东一行到访奇安信

P46 | 首批、首家！奇安信椒图通过可信安全“云上勒索攻击防护”检验

P47 | 奇安信获评“2023 年度漏洞报送最具贡献单位”

P48 | 连续三年排名第一 奇安信领跑 CWPP 云工作负载保护平台市场

P49 | 奇安信四项信创解决方案全部上榜“2024 广东软件风云榜”

P50 | 《云安全资源池能力指南》发布 奇安信荣获创新力及市场执行力双第一

P51 | 三大市场持续领先！奇安信终端安全、数据安全、分析和情报市场再获第一

P52 | “眼明心安”项目荣获“520 社会责任日”关爱儿童议题优秀案例

P53 | “眼明心安”2024 公益筛查带教活动出征仪式在京举行

政策篇



国内，新技术安全备受关注。两部门发文开展“网络去 NAT”专项工作，要求加强 IPv6 安全防护管理，《国家人工智能产业综合标准化体系建设指南（2024 版）》印发，提出规范人工智能全生命周期的安全要求，针对智能网联汽车车外数据收集场景，全国网安标委发布规范指南征求意见；

国际上，美议员提出多项法案加强国家网络安全。就破坏性网络攻击频发现象，美议员提案评估关基设施切换手动操作的可行性，就联邦监管“九龙治水”现象，美议员提案成立协调机构识别改进“繁冗、不一致或互相矛盾”的网络安全法规。



国家标准《数据安全技术 个人信息保护合规审计要求》公开征求意见

7 月 12 日，全国网络安全标准化技术委员会归口的国家标准《数据安全技术 个人信息保护合规审计要求》已形成标准征求意见稿，现公开征求意见。该文件提出了个人信息保护合规审计原则，规定了个人信息保护合规审计的实施要求，适用于个人信息处理者开展个人信息保护合规审计工作，也可对相关机构对个人信息处理活动进行个人信息保护合规审计提供参考。



两部门印发《关于开展“网络去 NAT”专项工作 进一步深化 IPv6 部署应用的通知》

7 月 10 日，工业和信息化部、中央网信办联合印发《关于开展“网络去 NAT”专项工作 进一步深化 IPv6 部署应用的通知》。该文件包括有序实现网络升级、持续拓宽 IPv6 通路、确保网络安全稳定等五方面工作任务。在网络安全方面，该文件要求基础电信企业、互联网企业针对开展 IPv6 升级改造的网络系统，要持续强化网络安全防护管理，定期开展风险评估，做好重要系统网络安全风险监测预警和应急处置；对于影响范围较大的设备升级、配置变更、服务变动，要制定预案，加强关键节点性能监测，保障网络安全稳定运行。



四部门联合印发《国家人工智能产业综合标准化体系建设指南（2024 版）》

7 月 2 日，工业和信息化部、中央网信办、国家发展改革委、国家标准委等四部门近日联合印发《国家人工智能产业综合标准化体系建设指南（2024 版）》。该文件提出，人工智能标准体系结构包括基础共性、基础支撑、关键技术、智能产品与服务、赋能新型工业化、行业应用、安全 / 治理等 7 个部分。在安全标准方面，要求规范人工智能技术、产品、系统、应用、服务等全生命周期的安全要求，包括基础安全，数据、算法和模型安全，网络、技术和系统安全，安全管理和服务，安全测试评估，安全标注，内容标识，产品和应用安全等标准。该文件还提出，到 2026 年我国新制定国家标准和行业标准 50 项以上，引领人工智能产业高质量发展的标准体系加快形成。



《网络安全标准实践指南——一键停止收集车外数据指引》公开征求意见

6 月 24 日，全国网络安全标准化技术委员会秘书处组织编制了《网络安全标准实践指南——一键停止收集车外数据指引（征求意见稿）》，现公开征求意见。该文件给出了在装有车载摄像头、雷达等传感器的汽车上设置一键停止收集车外数据功能的指引，适用于汽车制造企业及相关零部件或服务提供商设计、开发、制造具有车外数据收集功能的汽

车，不适用于主驾无人的高级别自动驾驶汽车。该文件旨在减少由于车载摄像头和雷达设备持续开启造成的重要数据和敏感个人信息违规收集、处理等安全问题，使汽车的辅助驾驶功能更好地应用。



国家标准《网络安全技术 关键信息基础设施安全保护能力指标体系》公开征求意见

6月20日，全国网络安全标准化技术委员会归口的《网络安全技术 关键信息基础设施安全保护能力指标体系》国家标准现已形成标准征求意见稿，公开征求意见。该文件规定了关键信息基础设施安全保护能力指标体系，包括基本保护级、强化保护级、战略保护级的安全保护能力指标，适用于指导关键信息基础设施运营者对关键信息基础设施安全保护能力的建设，也可保护工作部门、国家监管部门及第三方评估机构等提供参考。



美议员提出《联邦网络安全法规简化法案》，以改进监管碎片化问题

7月8日，美国参议员 Gary Peters 和 James Lankford 共同提出《联邦网络安全法规简化法案》，要求白宫国家网络总监成立一个跨机构委员会，负责协调联邦监管机构对企业提出的各种网络安全要求。该法案提出一个月前，参议院举行了一场听证会。听证会期间，国家网络总监办公室负责网络政策和项目的助理 Nicholas Leiserson 警告议员，网络安全法规的“碎片化”问题日益严重。他说：“为了解决这个问题，需要国家网络总监办公室和国会牵头，由私营部门提供信息。”该法案解决了这次听证会上提出的许多问题。跨机构委员会将负责识别“过于繁冗、不一致或相互矛盾”的网络安全要求，并提出更新建议，同时在各机构之间建立最低标准和互惠机制。



美国空军部发布《空军部零信任战略》

7月2日，美国空军部首席信息官办公室发布《空军部零信任战略》，旨在加强空军部网络安全态势，为作战人员提供有保障且安全的数据访问，以应对战争并阻止对手实现信息优势。该文件结合《国防部零信任战略》提出了7个战略目标，包括实现应用级可见性和控制、数据成为新的边界、向合理的实体和诉求提供合理的访问权限、降低单一设备风险、随时随地访问受保护资源、基于安全策略的自动安全响应、改进检测和响应时间。



美国国防部发布《支点：国防部信息技术发展战略》

6月25日，美国国防部发布《支点：国防部信息技术推进战略》，旨在利用技术力量推动变革并催化作战人员数字化现代化，为更好地协调信息技术运用以推进美国国防部优先事项提供路线图。该战略提出了四条工作方向，以确保美国国防部能够继续为国家的作战人员及其支持人员提供联系、保护并执行任务，具体包括提供联合作战信息技术能力、实现信息网络与计算现代化、优化信息技术治理、培养一支卓越的数字化人才队伍。



美国众议院议员提出《关键基础设施应急计划法案》

6月18日，美国众议院常设情报委员会成员 Dan Crenshaw、众议院国土安全委员会成员 Seth Magaziner 共同提出《关键基础设施应急计划法案》，要求美国网络安全和基础设施安全局局长与联邦紧急事务管理局局长及其他部门风险管理机构合作，向国会提交一份部门间联合评估报告，评估网络攻击期间关键基础设施的手动操作。具体评估内容包括关键基础设施无法迅速切换手动操作时国家网络事件响应计划的预案、CISA 支持关键基础设施事件响应的能力与义务、FEMA 国家响应框架对关键基础设施的支持情况、关键基础设施切换手动操作面临的潜在成本和挑战等。



全球关基设施网络威胁态势不断恶化。印尼国家数据中心遭勒索攻击，导致边检等服务中断多天，超 210 个政府机构被波及；菲律宾国家医保系统泄露超 4200 万用户数据，负责管理该系统的国企高管遭议会公开质询；猴痘病毒疫情爆发期间，南非国家卫生实验室因勒索攻击中断服务。



南非矿业巨头遭网络攻击：被迫隔离 IT 系统 企业运营受干扰

7 月 11 日 MyBroadband 消息，全球最大的铂金和黄金生产商之一、南非矿业巨头 Sibanye-Stillwater 透露，其全球 IT 系统遭受了网络攻击。该公司向利益相关者发出通知称：“公司发现事件后，便立即根据事件响应计划实施了主动隔离 IT 系统、保护数据的应急措施……虽然对事件的调查仍在进行中，但可以肯定事件对集团全球运营的干扰较有限。”该公司进一步表示：“我们自愿向相关监管机构报告此事件，并将在必要时提供进一步更新。”



美国大型银行因勒索攻击关闭系统多天，客户无法访问账户或交易

7 月 2 日 Bleeping Computer 消息，美国最大的信用合作社之一 Patelco 披露，公司 6 月 29 日遭遇了一次勒索软件攻击，为了控制事件影响主动关闭了多个面向客户的银行系统。Patelco 公告显示，网银、移动应用、电子交易等多类服务均不可用，借记卡、信用卡交易服务可用单功能受影响。7 月 3 日，Patelco 透露核心系统已经通过网络安全专家检查，客户资金安全不存在问题。由于此次勒索软件攻击事件，超 40 万客户至少连续 5 天无法访问账户或进行交易。



菲律宾国家医保系统泄露超 4200 万用户数据，负责国企高管遭议会公开质询

7 月 9 日 The Record 消息，菲律宾健康保险公司

(PhilHealth) 执行副总裁 Eli Dino Santos 日前在菲律宾众议院拨款委员会听证会上作证称，该公司发生违规事件后尚未按照法律要求，向每位受害者发送有关数据泄露的通知，此举遭到议员们的猛烈抨击。菲律宾健康保险公司是由该国卫生部长担任董事长的国企，负责管理国家医保系统，该系统在 2023 年 9 月遭遇勒索软件攻击，导致服务中断数周，且超 4200 万用户数据泄露。按照法律要求，菲律宾健康保险公司需要在事件发生后 72 小时内通知受害者，但其未能执行。



著名远控软件 TeamViewer IT 系统遭 APT 攻陷，安全专家建议暂时删除

6 月 29 日 The Record 消息，国际知名远程连接软件厂商 TeamViewer 在 28 日确认，一家极其活跃的俄罗斯黑客组织在本周早些时候入侵了其内部 IT 环境。TeamViewer 后续将此次攻击事件归咎于 APT29，据悉该组织隶属于俄罗斯的对外情报局 (SVR)，曾发起 2020 年的 SolarWinds 黑客事件和 2016 年对美国民主党全国委员会的攻击。TeamViewer 公司多次更新公告强调，此次攻击的受影响范围仅限内部 IT 环境，不涉及产品、TeamViewer 连接平台或任何客户数据。由于近年来供应链攻击屡创纪录，有安全专家建议暂时删除 TeamViewer。



印尼国家数据中心遭勒索攻击：边检等服务中断数天 超 210 个政府机构被波及

6 月 24 日路透社消息，印度尼西亚通信部长透露，一名网络攻击者入侵了该国国家数据中心，干扰了机场的边检

工作，导致出入境柜台排起了长队。据悉，位于泗水的印尼国家数据中心在上周遭到 Lockbit 勒索软件变种攻击，导致超过 210 家中央和地方政府机构受到波及，其中以移民边检服务受影响最为严重，中断了约 3 天时间，攻击者向其索要 800 万美元赎金。印尼通信部官员 Samuel Abrijani Pangerapan 表示，数字取证调查正在进行中，目前尚未找到更多细节。



猴痘病毒爆发期，南非国家卫生实验室因勒索攻击中断服务

6 月 24 日 MyBroadband 消息，南非国家卫生实验室服务局（NHLS）上周末遭到入侵，已关闭 IT 系统。该部门的电子邮件、网站及检索和存储患者实验室测试结果的系统均已离线。NHLS 首席执行官 Koleka Mlisana 教授发布备忘录，表示入侵造成了损害，说明 NHLS 遭受了勒索软件感染或类似的破坏性攻击。为应对此次攻击，NHLS 已经实施了“停机协议”，以确保提供必要的资源来解决这种情况。当前南非正处于猴痘病毒爆发期，NHLS 此前积压了大量待检测样本。此次攻击虽然未影响 NHLS 的内部运转，但 IT 中断导致其无法和外界进行数据通信，只能依赖电话或打印机来传递紧急测试结果。



供应商被黑导致美国上万家汽车经销商业务停顿

6 月 19 日 Bleeping Computer 消息，美国汽车经销商软件 SaaS 厂商 CDK Global 日前遭受大规模网络攻击，导致公司关闭系统，大量客户无法正常经营业务。从 6 月 19 日夜间到 20 日白天，CDK Global 遭遇网络攻击。为防止攻击蔓延，该公司被迫关闭 IT 系统、电话和应用程序。汽车经销商 IT 安全和服务公司 Proton Dealership IT 的首席执行官 Brad Holton 表示，攻击导致 CDK 在 20 日凌晨两点左右下线了两个数据中心。由于 CDK 公司被黑后关闭了业务系统，下游的汽车经销商们业务遭受广泛中断，包括跟踪和订购汽车零部件、销售新车、提供融资等。多位知情人士称，此次攻击系 BlackSuit 勒索软件团伙所为，据悉该组织前身为 Royal 勒索软件团伙。



全球叉车巨头科朗遭网络攻击，工厂生产被迫中断超一周

6 月 19 日 Bleeping Computer 消息，美国叉车制造巨头科朗设备（Crown Equipment）确认，本月早些时候遭受了一次网络攻击，导致其工厂的制造活动中断。大约 6 月 8 日起，科朗员工陆续报告公司遭遇了安全入侵，所有 IT 系统关闭。员工被告知不要接受多因素认证（MFA）请求，并对网络钓鱼邮件保持警惕。由于 IT 系统关闭，员工无法记录工时、访问服务手册，在某些情况下也无法交付机械设备。据悉，这次入侵是因为一名科朗员工遭到社交工程攻击，使得威胁行为者在他的电脑上安装远程访问软件。



AMD、苹果疑遭网络攻击，产品源代码泄露

6 月 18 日 Bleeping Computer 消息，AMD 公司疑似发生了“源代码级别”的数据泄露，该公司正在调查是否遭遇了一次网络攻击。黑客组织 IntelBroker 在地下论坛上出售声称从 AMD 窃取的数据，包含 AMD 员工信息、财务文件及源代码等机密信息。IntelBroker 分享了一些据称被盗的 AMD 凭证截图，但尚未透露这些数据的售价或其获取方式。据 Dark Web Informer 消息，IntelBroker 还在地下论坛发布了个苹果内网工具的源代码，包括 AppleConnect-SSO、Apple-HWE-Confluence-Advanced、AppleMacroPlugin。IntelBroker 声称这些数据是苹果公司 6 月泄露的。



勒索攻击致使英国首都近千台手术被迫取消

6 月 14 日 Bleeping Computer 消息，英国国家医疗服务体系（NHS）近日透露，6 月初 Synnovis 医疗组织遭到勒索软件攻击，导致伦敦多家医院受到影响，被迫取消了数百项手术计划和门诊预约。据悉，俄罗斯勒索软件组织 Qilin 发起攻击，导致 Synnovis 系统被锁定，盖伊和圣托马斯 NHS 信托基金会、国王学院医院 NHS 信托基金会，以及整个伦敦东南部的初级保健提供商都遭遇持续服务中断。各家医院发布的官方备忘录显示，这一“持续的重大事件”对医疗程序、手术、输血和血液检测等造成了“重大影响”。



漏洞篇



7月初，OpenSSH 远程命令执行漏洞 CVE-2024-6387 技术细节和 PoC 公开引发舆论恐慌，由于该软件使用范围极广，有观点认为这是又一起严重漏洞事故。据奇安信 CERT 分析，当前的漏洞利用代码利用过程复杂、成功率不高且耗时较长，利用可能性为中级，鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



ServiceNow 多个高危漏洞安全风险通告

7月11日，奇安信 CERT 监测到官方修复 ServiceNow UI Jelly 模板注入漏洞 (CVE-2024-4879) 和 ServiceNow Glide 表达式注入漏洞 (CVE-2024-5217)。ServiceNow 的 Jelly 模板和 Glide 表达式由于输入验证不严格，存在注入漏洞。这些漏洞可以被未经身份验证的攻击者通过构造恶意请求利用，在 ServiceNow 中远程执行代码。目前该漏洞技术细节与 PoC 已在互联网上公开，鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



GitLab 身份认证绕过漏洞安全风险通告

7月11日，奇安信 CERT 监测到官方修复 GitLab 身份认证绕过漏洞 (CVE-2024-6385)，攻击者可以在某些情况下以其他用户的身份触发 pipeline，从而造成身份验证绕过。奇安信鹰图资产测绘平台数据显示，该漏洞关联的国内风险资产总数为 48,6206 个，关联 IP 总数为 24,776 个。鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



Splunk Enterprise 任意文件读取漏洞安全风险通告

7月8日，奇安信 CERT 监测到 Splunk Enterprise 任意文件读取漏洞 (CVE-2024-36991) 在互联网上公开，在特定版本的 Windows 版 Splunk Enterprise 中，未经身

份认证的攻击者可以在 /modules/messaging/ 接口通过目录穿越的方式读取任意文件，造成用户信息的泄露。奇安信鹰图资产测绘平台数据显示，该漏洞关联的国内风险资产总数为 2374 个，关联 IP 总数为 501 个。目前该漏洞技术细节与 EXP 已在互联网上公开，鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



Rejetto HTTP File Server 模板注入漏洞安全风险通告

7月8日，奇安信 CERT 监测到 Rejetto HTTP File Server 模板注入漏洞 (CVE-2024-23692)，由于该系统存在模板注入漏洞，未经身份验证的攻击者通过发送特制的 HTTP 请求在受影响的系统上执行任意命令。奇安信鹰图资产测绘平台数据显示，该漏洞关联的国内风险资产总数为 84,919 个，关联 IP 总数为 13,059 个。目前该漏洞已发现在野利用，鉴于该漏洞 EXP 已广泛传播，利用成本极低，请您立即着手修复，以避免潜在风险。



Apache Tomcat 拒绝服务漏洞安全风险通告

7月5日，奇安信 CERT 监测到官方修复 Apache Tomcat HTTP/2 拒绝服务漏洞 (CVE-2024-34750)，该漏洞是一个拒绝服务漏洞，Apache Tomcat 在处理 HTTP/2 流时，未能正确处理某些异常 HTTP 头情况，导致 HTTP/2 流计数错误，从而导致处理该请求时允许无限超时，无法关闭本应该终止的连接。鉴于此漏洞影响范围较大，建议客户尽快做好自查及防护。



Artifex Ghostscript 格式字符串漏洞安全风险通告

7月3日，奇安信 CERT 监测到官方修复 Artifex Ghostscript 格式字符串漏洞 (CVE-2024-29510)，Ghostscript 中存在格式字符串注入可导致 Shell 命令执行，未经身份认证的本地攻击者通过特制文件利用该漏洞可以实现代码执行，在 Linux 操作系统上可用于鱼叉攻击，如将恶意 esp 文件嵌入到 LibreOffice 文档文件，用户通过 LibreOffice 打开文档将导致代码执行。目前该漏洞技术细节与 EXP 已在互联网上公开，鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



GeoServer 远程代码执行漏洞安全风险通告

7月3日，奇安信 CERT 监测到官方修复 GeoServer 远程代码执行漏洞 (CVE-2024-36401)，由于该系统不安全地将属性名称解析为 XPath 表达式，未经身份认证的远程攻击者可以通过该漏洞在服务器上执行任意代码，从而获取服务器权限。目前该漏洞技术细节与 EXP 已在互联网上公开，鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



OpenSSH 远程代码执行漏洞安全风险通告

7月1日，奇安信 CERT 监测到官方修复 OpenSSH 远程代码执行漏洞 (CVE-2024-6387)，该漏洞是由于 OpenSSH 服务器 (sshd) 中的信号处理程序竞争问题，未经身份验证的攻击者可以利用此漏洞在 Linux 系统上以 root 身份执行任意代码。目前该漏洞技术细节和 PoC 已在互联网上公开，该漏洞为之前 CVE-2006-5051 的二次引入，当前的漏洞利用代码仅针对在 32 位 Linux 系统上运行的 OpenSSH，64 位 Linux 系统上利用该漏洞的难度会更大，在 Linux 系统上以 Glibc 编译的 OpenSSH 上成功利用，不过利用过程复杂、成功率不高且耗时较长。平均要大于 10,000 次才能赢得竞争条件，需要 6~8 小时才能获得远程 root shell。在以非 Glibc 编译的 OpenSSH 上利用此漏洞也是可能的，但尚未证实。虽然目前还没有发现真正实现远程代码执行的 PoC，鉴于该漏洞影响范围较大，建议客

户尽快做好自查及防护。



GitLab 身份认证绕过漏洞安全风险通告

6月28日，奇安信 CERT 监测到官方修复 GitLab 身份认证绕过漏洞 (CVE-2024-5655)，攻击者可以在某些情况下以其他用户的身份触发 pipeline，从而造成身份验证绕过。奇安信鹰图资产测绘平台数据显示，CVE-2024-5655 关联的国内风险资产总数为 1,595,223 个，关联 IP 总数为 112,251 个。鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。



Progress MOVEit Transfer 身份认证绕过漏洞安全风险通告

6月26日，奇安信 CERT 监测到 Progress MOVEit Transfer 身份认证绕过漏洞 (CVE-2024-5806) 在互联网上公开，该漏洞是一个严重的认证绕过缺陷，它使得攻击者能够在没有任何有效凭证的情况下，通过特定的技术手段来冒充服务器上的任何用户，从而获取对敏感数据的未授权访问。奇安信鹰图资产测绘平台数据显示，CVE-2024-5806 关联的国内风险资产总数为 6637 个，关联 IP 总数为 227 个。目前该漏洞技术细节与 PoC 已在互联网上公开，鉴于此漏洞影响范围较大，建议客户尽快做好自查及防护。



Ollama 远程代码执行漏洞安全风险通告

6月25日，奇安信 CERT 监测到官方修复 Ollama 远程代码执行漏洞 (CVE-2024-37032)，在没有强制身份验证的 Ollama 服务器上，攻击者可通过操控服务器接口下载恶意文件，从而导致远程代码执行。奇安信鹰图资产测绘平台数据显示，CVE-2024-37032 关联的国内风险资产总数为 3955 个，关联 IP 总数为 994 个。目前该漏洞技术细节与 EXP 已在互联网上公开，鉴于该漏洞影响范围较大，建议客户尽快做好自查及防护。

注：使用公司邮箱发送企业名称和需开通订阅的邮箱地址至 cert@qianxin.com，即可申请订阅最新漏洞通告。

规划
快一步

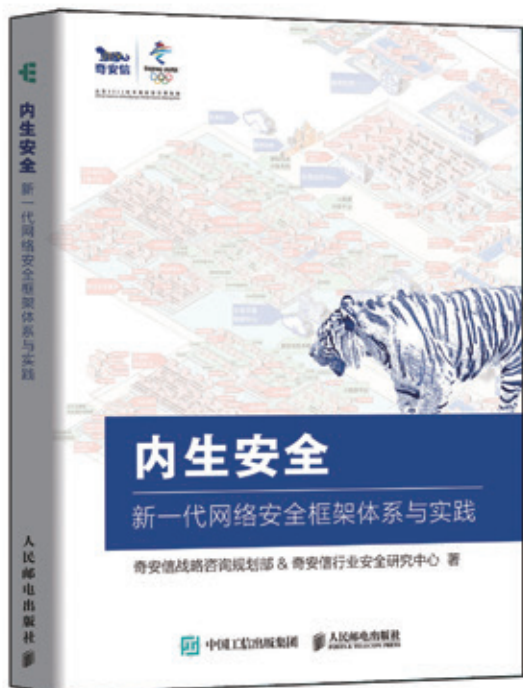


北京2022年冬奥会官方赞助商
Official Sponsor of the Olympic Winter Games Beijing 2022

新书发布

内生安全权威解读

19支团队、37位专家倾力打造
政企“十四五”网络安全规划必读书籍



什么是内生安全

内生安全从何而来

为什么要内生安全

内生安全如何落地

新一代网络安全框架

“十工五任”建设要点

扫描二维码
专享内购价



前沿

奇安信的AI

驱动威胁情报运营实践

作为网络安全领域的重要变革力量，人工智能正在重塑组织的威胁情报运营方式。作为国内网络安全的领军企业，奇安信走在 AI 驱动威胁情报运营的创新前沿，正在利用人工智能和大模型的力量，为客户提供所需的威胁情报，以更好地保护资产并领先威胁者一步。



前沿应用： 奇安信 AI 驱动威胁情报运营实践

本文撰稿 汪列军

作为网络安全领域的重要变革力量，人工智能正在重塑组织的威胁情报运营方式：利用人工智能可以自动获取威胁情报，帮助安全人员全面了解安全态势，做出更明智的安全决策。

作为国内网络安全的领军企业，奇安信走在 AI 驱动威胁情报运营的创新前沿，正在利用人工智能和大模型的力量为客户提供所需的威胁情报，以更好地保护资产并领先威胁者一步。

一、威胁情报运营各环节及 AI 技术的应用概览

下面简单介绍奇安信威胁情报运营的基本流程，这本质上是数据驱动的过程（见图1）：从多个源头做数据采集，经过一系列清洗、分类、检测、关联、结构化等过程处理以后，形成数据或服务输出给安全检测产品和相应的威胁分析响应人员。



图 1 奇安信威胁情报运营的基本流程

在网络威胁情报的运营生产过程中，包括大模型在内的AI技术在流程的各个环节加以合理运用，能极大提升运营的效率与效果。



图 2 奇安信威胁情报运营主要环节及 AI 技术应用情况

图 2 展示了奇安信威胁情报运营的主要环节及相应 AI 技术应用情况,这是多年累积的结果。传统的机器学习的技术早已得到应用,在大模型技术出现以后有些能力被取代或加强。

首先，从数据采集开始。在这个阶段，我们的目标是从各种来源收集尽可能多的数据，并进行快速而准确地分类和提取。这里需要执行网页解析、文本翻译和 OCR 识别等操作。几乎在获取的同时对信息进行分类，去除噪声和无关数据。运用自然语言处理技术，特别是利用大模型进行语义标签的自动分类和智能分析，从而完成精确的多维度的信息分类，据此可以进入不同的处理子流程。其中，大语言模型发挥其在数据分析方面的核心优势。它们能够理解和处理大量非结构化的文本及音视频数据，并从中提取我们需要的关键信息。

然后，是检测识别。这一阶段的目标是对文件样本与网络流量等元数据执行检测，进行威胁判定，标记恶意实体，这是威胁情报运营过程中的核心环节，安全厂商的能力所在。本阶段使用了比较成熟的机器学习算法和深度学习技术，实现恶意样本的标定和恶意网络攻击的识别，收集其相关的威胁情报工件数据。

接下来，我们会对收集到的**威胁元数据进一步结构化**，综合利用自然语言处理技术、大模型技术及大规模威胁知识图谱的构建，提取精准的威胁实体，并建立多类型的关联。

在结构化之后，我们进入**信息富化的阶段**。这一步骤涉及关键子图的提取和大规模图嵌入，从其他源补充更多的关联信息。通过这些方法，我们可以揭示出隐藏在数据背后的模式和联系。

实体信息富化以后做**进一步的拓线关联**。我们利用智能 Agent 和图相似计算来推荐可能的相关恶意实体和事件。这使得我们能够建立一个完整的威胁图景，并挖掘出更多潜在的威胁对象实体。

最后，我们生成报告并对情报进行摘要总结。这一步依赖于大语言模型的技术，它们能够生成简洁且易于理解的报告并给出结论，基于公开和内部信息构建安全知识库，为威胁分析和应急响应人员提供信息支持。

这就是整个网络威胁情报生产及运营的过程。通过结合大语言模型、安全大数据计算技术和深度学习技术的综合应用，实现威胁情报的快速发现和输出。

在测试和应用 AI 技术的过程中，我们用到了一些工具。底层的当然是一些基础大模型，比如，最强大的来自 OpenAI 公司的闭源模型 GPT4，国内比较流行的口碑较好、测试下来效果也确实不错的 Kimi 和通义千问，还有能力与 GPT4 可以媲美的开源模型 LLaMA3，Mistral 也是一个性能与能力平衡得较好的模型，当然包括奇安信自研的安全大模型 QAX-GPT，安全相关的知识与能力更优。网络威胁情报生产运营过程及使用工具见图 3。



图 3 网络威胁情报生产运营过程及使用工具

基础模型的能力通过 Web 界面和 API 可以直接被使

用，做文本分类、模式识别和信息提取。在此之上，有各种应用程序，数据获取，如 ScrapeGraphAI 和 Jina；各类支持开发框架，如 Ollama、LlamaIndex、LangChain、DIFY；组合起来开发实现很强功能的智能体应用，目前非常好用的搜索增强平台，如 Perplexity 及国内的秘塔；各类本地 RAG 工具，如 QAnything、AnythingLLM 等。可以说，整个基于 AI 大模型的生态已经快速建立起来了。

二、情报运营各环节的 AI 技术测试应用及体会

1、网页解析

威胁情报的运营开始于信息的收集，来自各种 Web 页面的开源内容是收集的主要目标。

利用大模型极其强大的文本分析能力，可以方便地实现网页的处理，ScrapeGraphAI 工具就是这样一个集成了大模型文本分析能力的网页内容爬取工具。下面是一个使用其爬取 Mitre 网站的 APT 组织信息页面并转化为结构化数据的例子，见图 4。

Groups	Name	Associated Groups	Description
APT1	APT1		
APT2	APT2		
APT3	APT3		
APT4	APT4		
APT5	APT5		
APT6	APT6		
APT7	APT7		
APT8	APT8		
APT9	APT9		
APT10	APT10		
APT11	APT11		
APT12	APT12		
APT13	APT13		
APT14	APT14		
APT15	APT15		
APT16	APT16		
APT17	APT17		
APT18	APT18		
APT19	APT19		
APT20	APT20		
APT21	APT21		
APT22	APT22		
APT23	APT23		
APT24	APT24		
APT25	APT25		
APT26	APT26		
APT27	APT27		
APT28	APT28		
APT29	APT29		
APT30	APT30		
APT31	APT31		
APT32	APT32		
APT33	APT33		
APT34	APT34		
APT35	APT35		
APT36	APT36		
APT37	APT37		
APT38	APT38		
APT39	APT39		
APT40	APT40		
APT41	APT41		
APT42	APT42		
APT43	APT43		
APT44	APT44		
APT45	APT45		
APT46	APT46		
APT47	APT47		
APT48	APT48		
APT49	APT49		
APT50	APT50		
APT51	APT51		
APT52	APT52		
APT53	APT53		
APT54	APT54		
APT55	APT55		
APT56	APT56		
APT57	APT57		
APT58	APT58		
APT59	APT59		
APT60	APT60		
APT61	APT61		
APT62	APT62		
APT63	APT63		
APT64	APT64		
APT65	APT65		
APT66	APT66		
APT67	APT67		
APT68	APT68		
APT69	APT69		
APT70	APT70		
APT71	APT71		
APT72	APT72		
APT73	APT73		
APT74	APT74		
APT75	APT75		
APT76	APT76		
APT77	APT77		
APT78	APT78		
APT79	APT79		
APT80	APT80		
APT81	APT81		
APT82	APT82		
APT83	APT83		
APT84	APT84		
APT85	APT85		
APT86	APT86		
APT87	APT87		
APT88	APT88		
APT89	APT89		
APT90	APT90		
APT91	APT91		
APT92	APT92		
APT93	APT93		
APT94	APT94		
APT95	APT95		
APT96	APT96		
APT97	APT97		
APT98	APT98		
APT99	APT99		
APT100	APT100		

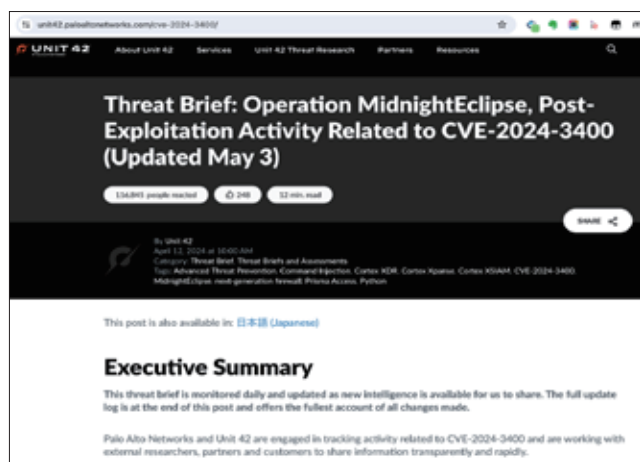
图 4 使用 ScrapeGraphAI 工具爬取 Mitre 网站 APT 组织信息

这一工具充分利用了大模型理解抽象指令和页面解析的能力，处理复杂页面不需要再专门写适配的解析代码，只要描述要求即可。当前这个例子核心 AI 能力使用了本地部署的 LLaMA3 模型，其实就当前的任务而言都不需要能力这么强的模型，用更轻量级的 Mistral 就可以了。

另一个非常有用的工具是 jina.ai，这是一个非常方便的

把网页内容精减成 Markdown 格式文本的工具，目前可以免费使用，使用频率上有点限制。它会获取网页，自动去除掉包括广告在内的无关信息，输出一个 Markdown 格式的文档，这些内容可以被后续的大模型应用所消费。

比如，给 jina.ai 指定下面这样的一个威胁行为体的介绍网页：



<https://unit42.paloaltonetworks.com/cve-2024-3400/>

jina.ai 会输出如下的 Markdown 文档：



使用方式非常简单，只需在要处理的网页链接前添加“https://r.jina.ai/”即可。

小结

• 利用大模型的能力进行网页爬取、解析及结构化目前还处于概念验证阶段，实用系统需要较强的工程化能力和强大的算力支持

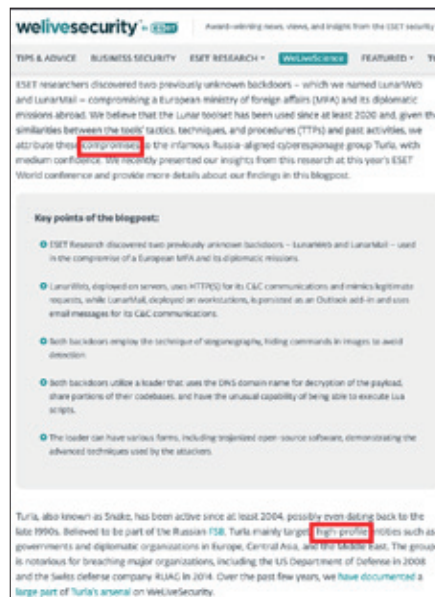
- 用于处理复杂网页较为合适，简单网页杀鸡用牛刀
- 获取页面内容最好用 jina 这样的核心信息提取服务，后续大模型处理时既可以减少非相关信息的干扰，又能减少 Token 的消耗
- 输出后续用于大模型的精减格式化数据非常有用，数据通过 API 以结构化的方式交换加速成为主流，将来 Web 页面被人访问的比例会越来越低，Web 可能慢慢走向消亡

2、文本翻译

我们获取到的文本可能是多语言的，主要是英文，需要尽可能准确地转换为中文，毕竟我们是中国的安全厂商，接下来看看大模型对于翻译问题最终解决。

下面的例子是安全厂商博客栏目发布的关于某个 APT 组织活动的介绍文章。我们重点关注两处可以体现翻译能力差距的地方。

- 1、Compromises
- 2、High-profile



<https://www.welivesecurity.com/en/eset-research/moon-backdoors-lunar-landing-diplomatic-missions/>

Compromises 最合适的翻译应该是失陷或侵入，high-profile 比较好的翻译应该高关注度、高知名度实体。

没有大模型加持的翻译工具，甚至连文本通顺都有问题，就象谷歌翻译（Google Translate）那样。

GPT4 翻译得比较通顺，但对我们关注的概念处理上还谈不上完全正确。

ESET的研究人员发现了两个之前未知的后门程序——我们将它们命名为LunarWeb和LunarMail——它们影响了一家欧洲外交部及其海外的外交使团。我们认为，自2020年以来至少已经在使用Lunar工具集，并且这些工具的技术、技术和程序（TTPs）与过去的活动之间的相似性，我们中等程度的相信这些安全漏洞与臭名昭著的俄罗斯黑客组织Turla有关。我们最近在本年度的ESET World会议上分享了这项研究的发现，并在这篇博客文章中提供了更多关于我们发现的细节。

博客文章的关键点包括：

- ESET研究发现了两个以前未知的后门——LunarWeb和LunarMail——这些后门被用来攻击一家欧洲外交部及其外交使团。
- LunarWeb部署在服务器上，使用HTTP(S)进行其命令与控制（C&C）通信，并模仿合法请求，而LunarMail部署在工作站上，以Outlook插件的形式存在，并使用电子邮件消息进行其C&C通信。
- 这两个后门都采用了隐写技术，通过图像中隐藏命令来避免检测。
- 这两个后门都使用了一个加载器，该加载器使用DNS域名来解密有效载荷，它们共享代码库的部分，并具有执行Lua脚本的不同寻常的能力。
- 加载器可以有多种形式，包括被篡改的开源软件，这展示了攻击者使用的高级技术。

Turla，也被称为Snake，至少从2004年起就一直活跃，可能甚至可以追溯到1990年代末。被认为是俄罗斯联邦安全局（FSB）的一部分，Turla主要针对的是欧洲、中亚和中东的政府机构和外交使团。该组织因突破包括2008年美国国防部和2014年瑞士防务公司RUAG在内的关键组织而闻名。在过去几年里，我们在WeLiveSecurity上记录了Turla的大量武器库。

我们目前的调查始于检测到一个加载器在一个未识别的服务器上解密并运行外部文件中的有效载荷。这引导我们发现了一个之前未知的后门，我们将其命名为LunarWeb。随后，我们在一家欧洲外交部的外交机构中发现了一个类似的链条，LunarMail在那里被部署。值得注意的是，攻击者还包含了第二个后门——我们将其命名为LunarMail——它使用了一种不同的命令和控制（C&C）通信方法。

GPT4-1, Prompt: 翻译如下文本到中文

GPT4 改进提示词以后的表现有所改善，但 Kimi 看起来翻译得最贴切。

ESET的研究人员发现了两个以前未知的后门程序——我们将它们命名为LunarWeb和LunarMail——这些后门程序危害了一家欧洲外交部（MFA）及其海外的外交使团。我们相信，自2020年起至少已经使用了Lunar工具集，用于这些工具的技术、技术和程序（TTPs）与过去的活动之间的相似性，我们中等程度的相信这些安全事件与臭名昭著的俄罗斯黑客组织Turla有关。我们最近在本年度的ESET World会议上展示了这项研究的发现，并在这篇博客文章中提供了更多关于我们发现的详细信息。

博客文章的关键点：

- ESET研究发现了两个以前未知的后门——LunarWeb和LunarMail——它们被用于入侵一家欧洲外交部及其外交使团。
- LunarWeb部署在服务器上，使用HTTP(S)进行其C&C（命令和控制）通信，并模仿合法请求，而LunarMail部署在工作站上，以Outlook插件的形式持久化，并使用电子邮件消息进行其C&C通信。
- 这两个后门都采用了隐写技术，隐藏命令在图像中以避免检测。
- 这两个后门都使用了一个加载器，该加载器使用DNS域名来解密有效载荷，共享它们的代码库的部分，并且具有执行Lua脚本的不寻常能力。
- 加载器可以有多种形式，包括被篡改的开源软件，展示了攻击者使用的高级技术。Turla，也被称为Snake，自2004年起就一直活跃，甚至可能可以追溯到1990年代末。据信是俄罗斯联邦安全局（FSB）的一部分，Turla主要针对欧洲、中亚和中东的政府和外交机构。该组织因入侵包括美国国防部在2008年和瑞士防务公司RUAG在2014年的重大组织而臭名昭著。在过去几年中，我们在WeLiveSecurity上记录了Turla的大部分武器库。

我们目前的调查始于检测到一个加载器在一个未识别的服务器上解密并运行一个外部文件中的有效载荷，这引导我们发现了一个以前未知的后门，我们将其命名为LunarWeb。随后，我们在一家欧洲外交部的外交机构中检测到了一个类似的链条，部署了LunarWeb。值得注意的是，攻击者还包含了第二个后门——我们将其命名为LunarMail——它使用了一种不同的命令和控制（C&C）通信方法。

GPT4-2, Prompt: 在网络安全语境下，翻译如下文本到中文

ESET研究人员发现了两个以前未知的后门——我们将其命名为LunarWeb和LunarMail——它们入侵了一家欧洲外交部（MFA）及其驻外外交使团。我们认为Lunar工具集至少自2020年以来一直在使用，并且鉴于这些工具的技术、技术和程序（TTPs）与过去的活动之间的相似性，我们中等信心相信这些入侵归因于与俄罗斯结盟的臭名昭著的网络安全组织Turla。我们最近在本年度的ESET世界大会上展示了我们从这项研究中获得的意见，并在这篇博客文章中提供了更多关于我们发现的细节。

博文的关键点：

- ESET研究发现了两个以前未知的后门——LunarWeb和LunarMail——用于入侵一家欧洲外交部及其外交使团。
- LunarWeb部署在服务器上，使用HTTP(S)进行其命令与控制（C&C）通信，并模仿合法请求，而LunarMail部署在工作站上，作为Outlook插件持久存在，并使用电子邮件消息进行其C&C通信。
- 两个后门都采用了隐写技术，将命令隐藏在图像中以避免被检测到。
- 两个后门都使用了一个加载器，该加载器使用DNS域名进行有效载荷的解密，共享了它们代码库的部分内容，并且具有执行Lua脚本的不寻常能力。
- 加载器可以有多种形式，包括特洛伊化的开源软件，展示了攻击者使用的高级技术。Turla，也被称为Snake，自2004年以来一直活跃，甚至可能可以追溯到1990年代末。人们认为它是俄罗斯联邦安全局（FSB）的一部分，Turla主要针对包括政府和外交机构在内的知名实体。这些实体位于欧洲、中亚和中东。该组织因入侵包括美国国防部在内的主要组织而臭名昭著，2008年入侵了美国国防部，2014年入侵了瑞士国防公司RUAG。在过去的几年中，我们在WeLiveSecurity上记录了Turla的大量武器库。我们当前的调查始于检测到一个加载器在未识别的服务器上解密并运行来自外部文件的有效载荷。这引导我们发现了我们命名为LunarWeb的以前未知的后门。随后，我们在一个欧洲外交部的外交机构中检测到了一个类似的链条，部署了LunarWeb。值得注意的是，攻击者还包含了第二个后门——我们将其命名为LunarMail——它使用了不同的命令和控制（C&C）通信方法。

Kimi, Prompt: 翻译如下文本到中文

小结

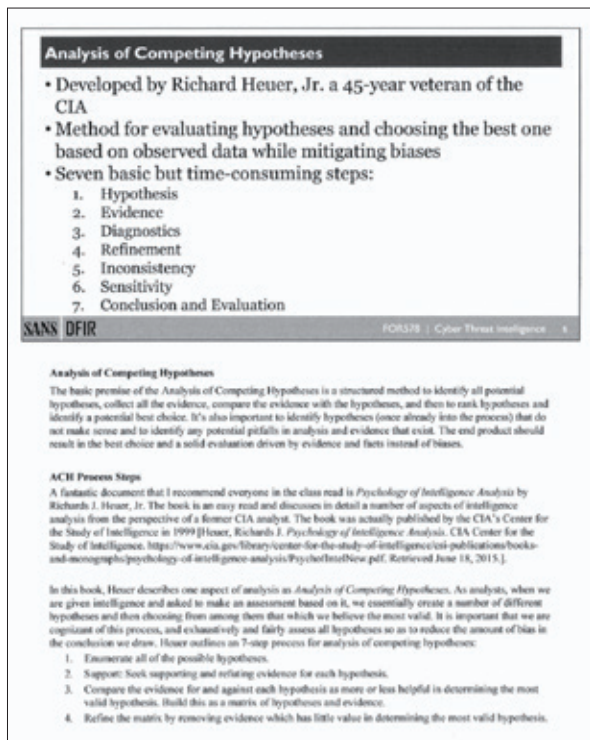
• 基于大模型技术的翻译水平碾压传统的翻译技术，大模型对多语言的掌握已经远超一般的人类，在技术文档的翻译方面代工基本上已经可有可无

• 就翻译中文的需求而言，中文的大模型翻译效果上优于哪怕是最好的国外大模型，在这点上国内原厂的大模型确实还有些优势

3、OCR 识别

情报信息的来源除了现成的文本，图像视频是另一个重要来源，在这个领域我们也应该积极获取数据，所以 OCR 的技术也是必备的，而这方面也是大模型所擅长的。

下面是一个扫描版的 SANS 培训材料，是 PDF 格式的文档，但文档的每页都是一张独立的图片。我们尝试通过 OCR 技术提取图片中的文本与格式。



SANS – 578.4 – Analysis and Dissemination of Intelligence (SANS).pdf

市面上已有不少基于深度学习的 OCR 工具，这里展示的 Surya 工具只是其中一个，对于文档的处理已经比较完善，识别完成以后保持与图像完全一致的观感（右边是源文档页图片，左边是识别后的文本）。

从情报分析需求方面，下面我们来看看大模型在 OCR 功能上的表现。

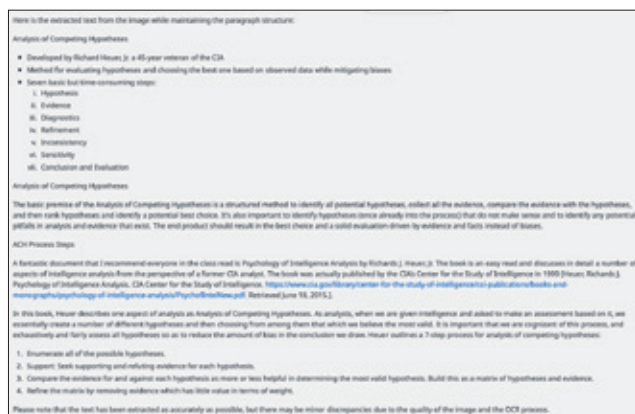


用大模型执行识别任务，先尝试 GPT4 vision 模型，

Prompt:

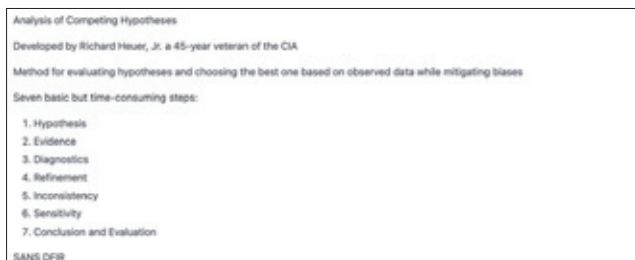
You are provided with one image (see uploaded file). You need directly extract text from image. You need feedback me the extracted text. Keep the paragraph structure from the original image. Remove extra new line characters and keep one new line between each paragraph.

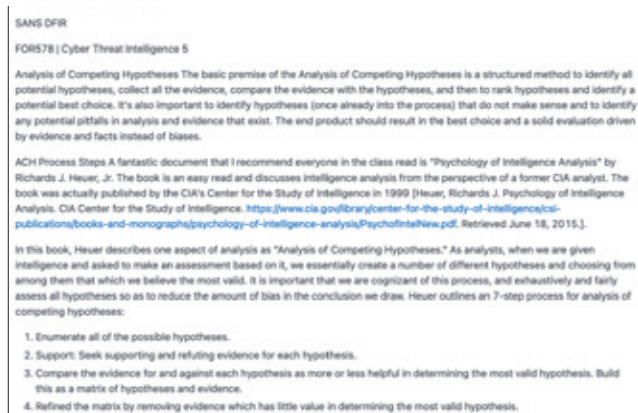
从识别以后输出的格式看，GPT4 并没有百分之百的还原，忽略了部分它认为不重要的信息，还会对格式有一些自己的改动，注意上面那 7 个条目的编号，PPT 页脚的内容被忽略了。



GPT4-vision 的识别结果

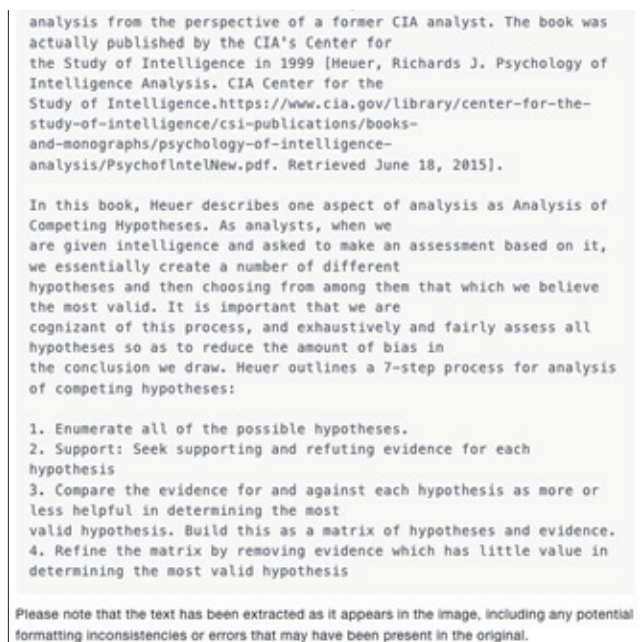
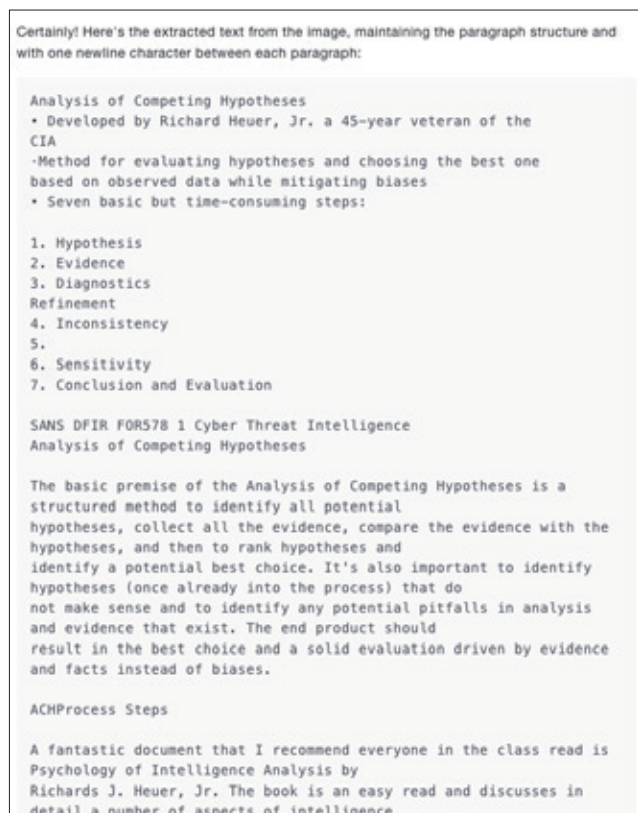
通义千问内容还原度相当高，没有什么自己额外的发挥，包括 PPT 页的页脚都翻译并呈现了，对段落进行了智能识别与重整。





通义千问的识别结果

Kimi 识别出了几乎所有文本内容，但在输出格式上做得很不好，甚至有点错乱，段落的回车位置却严格按照原图，没有做智能化的处理。



Kimi 的识别结果

总体来看，通义千问的输出效果最好。

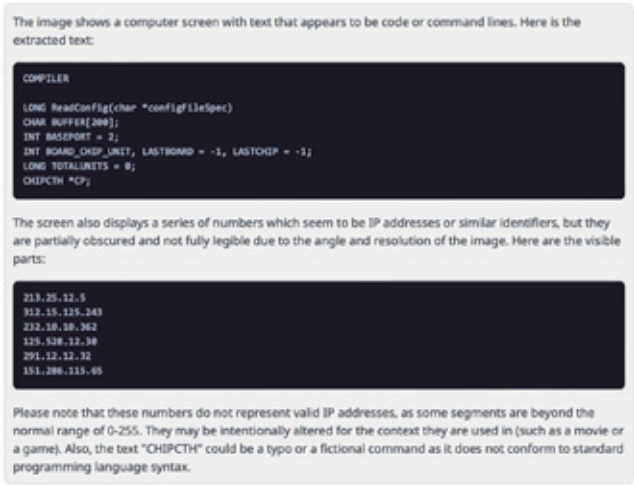
换到对于视频中的文字识别场景，我们先用 GPT4 vision 处理，Prompt:

You are provided with one image (see uploaded file). You need directly extract text from image. You need feedback me the extracted text.



电影《剑鱼行动》中的场景

可以看到，它准确地识别出了核心内容，还自动分割了区域，并给出了自己的看法。



GPT4 的识别输出

让 Kimi 和通义千问来识别一下同样这个图片，结果通义千问只识别出了部分内容，Kimi 识别出了里面大多数文字，但看起来并没有理解里面的内容。
但智能化的分析理解就一定好吗？看下面的例子。



这是某年参加的一个安全会议上拍的演讲题目，拍得太不清楚，让 GPT 来处理这个内容识别。



GPT4 的识别输出

注意 PPT 标题部分，出现了幻觉！大模型知识内置的好处在于能自动弥补上下文，甚至基于此做更深入的研判，但坏处是随意发挥造成误导。对于质量较差的图片，其输出的结果不能完全信任。
通义千问对这个图片的内容识别比 GPT4 准确得多：



2021年8月5日发现新上线域名及IP:
www.parwa.info
parwa.info
212.33.46.184

通过知道创宇NDR对某关键出口流量进行监测发现:

2021年8月11日,发现某组织使用 C2.independence day.pafwa.info发起攻击,改样本为:
md5: 4df2aef34ad8987c7f991d5ada7fd68

恶意文件名: Invitation Pass.pdf.Link
url:
https://independence day.pafwa.info/87/1/39/2/0/0/1
816006329/SiErdP3idkTAqhlIABLOQaqIHINw9Q
qubvSE/A/files/2e61563bhta

通义千问的识别输出

小结

- 对一般印刷文本的识别上,基于神经网络的 AI 技术已经能够非常准确地完成任务
- 由于大模型具备的理解能力,它不仅仅简单识别其中包含的文本,还能对场景做出判断,理解其中的意义,对文本的呈现会选择合适的方式,甚至利用模型内部的知识做补齐(不一定好事情)
- 涉及中文的处理最好还是用国内的大模型,如通义千问

大模型的智能可以是你的福音,但也可能成为你的毒药,严肃使用识别出来的信息需要进一步验证。

4、情报自动分类—传统神经网络方法

每天从数千个源收集到数以万计的文本以后,不可能由人工去挨个查看,需要有自动化的系统批量对文本做属性判定并在后续提取关键信息。

在大模型技术前,奇安信人工智能研究院就已经使用神经网络做文本分类,监督式的学习,核心用到了 Bert 模型。



之前的这个技术流程可以实现分类的粗分类,包括的属性有:是否安全相关、是否安全漏洞相关、是否恶意软件相关、是否网络犯罪相关、是否 APT 相关等,效果基本可以接受,能力已经加入了我们的威胁情报运营流程。



随着大模型出现,现在的技术都能分辨说话的情绪倾向,高兴还是沮丧,用来分类文本是否与安全相关更是轻而易举。这里是 Kimi 对一个网页内容研判的例子, Kimi 总结了网页的内容并准确地判定文章与网络安全无关。



Kimi 的分类判定

另一个介绍 Mac 平台恶意代码的文章, Kimi 非常轻松地对文章内容做的总结和判断。



1. 恶意软件发现：文章由Intego发布，该公司专注于Mac安全，并宣布发现了一种新的模仿Homebrew的“Cuckoo” Mac恶意软件。
2. 恶意软件行为：文章描述了这种恶意软件如何作为特洛伊木马传播，通常伪装成盗版或破解版本的应用程序，并利用恶意的Google广告进行传播。
3. 安全研究：文章提到了由安全研究员Alden在社交媒体上发布的关于这一新活动的信息，并提供了相关的博客文章链接。
4. 恶意软件分析：Intego的研究团队对恶意软件进行了分析，发现它具有在虚拟机环境中运行时改变行为的新功能，这使得恶意软件分析更加困难。
5. 安全建议：文章提供了如何保护Mac免受“Cuckoo”和其他恶意软件感染的建议，包括使用Intego VirusBarrier等防病毒软件。
6. 威胁指标：文章列出了与Homebrew相关的Cuckoo/AMOS恶意软件活动的指标（IOCs），包括恶意软件样本的SHA-256哈希值和相关域名及IP地址。
7. 安全意识提升：文章鼓励用户不要仅仅依赖搜索引擎找到合法网站，而应该养成使用书签的习惯，并提供了关于如何获取更多信息的资源链接。

文章内容涉及恶意软件的发现、分析、防护措施以及如何提高网络安全意识等多个方面，这些都是网络安全领域的重要组成部分。

Kimi 的分类和总结

但是，Kimi 对于某些具体的安全属性判定存在问题。文章安全漏洞应该是不相关的，与具体的网络犯罪也不相关，说明 Kimi 在这些问题的处理上存在缺陷，模型的能力并不完美。

那么我们来试一下其他的模型，Llama3 对于我们关心的几个问题是否判定是正确的，体现 Llama3 有更准确的内在知识。

请判断如下这个链接里的内容是否与网络安全有关，是否与具体的安全漏洞有关，是否与具体的恶意软件有关，是否与具体的数据泄露事件有关，是否与具体的网络犯罪活动有关，是否与具体的APT（Advanced Persistent Threat，高级威胁行为体）有关，是否包含具体的IOC（Indicator Of Compromise）实体，只要给出是与否的结论：<https://r.jina.ai/https://Offset.net/reverse-engineering/analyzing-com-mechanisms-in-malware/>

Here are the judgments:

1. 是否与网络安全有关: **YES**
2. 是否与具体的安全漏洞有关: **NO**
3. 是否与具体的恶意软件有关: **YES**
4. 是否与具体的数据泄露事件有关: **NO**
5. 是否与具体的网络犯罪活动有关: **NO**
6. 是否与具体的APT（Advanced Persistent Threat，高级威胁行为体）有关: **NO**
7. 是否包含具体的IOC（Indicator Of Compromise）实体: **NO**

The article appears to be a technical analysis of malware's communication mechanisms, which falls under the category of network security. It discusses reverse-engineering malware and analyzing its communication mechanisms, which is a topic related to malware analysis and cybersecurity.

Llama3 模型的属性判定

通义千问与 Llama3 的判断完全一致，而且正确，在这方面要比 Kimi 强。



通义千问的属性判定

小结

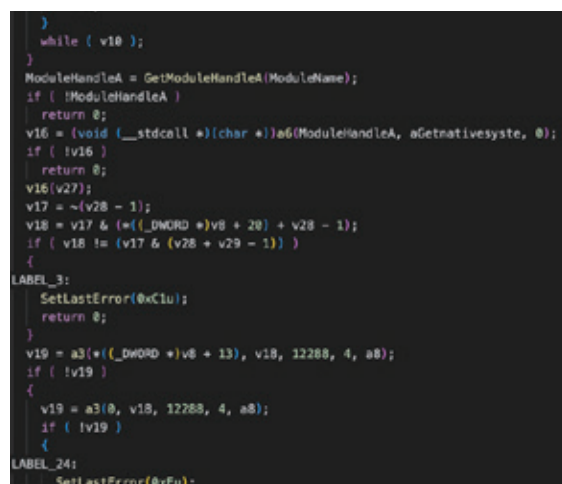
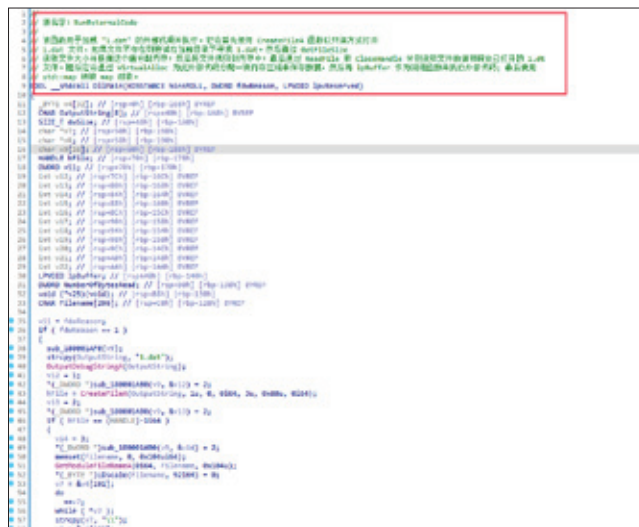
- 自然语言处理上，大模型对传统NLP形成碾压式优势，以前非常麻烦的流程设计实现，现在你需要提要求就行了
- 各个模型对情报分类及其他任务上表现还是有区别的，需要仔细评估擅长点，根据不同的任务选择合适的模型，资源丰富的话可以同时使用，采用投票机制
- 以上这些例子其实只是充分体现了大模型的在文本和图像处理上通用能力，跟具体的安全业务紧密度不太高，下面的任务就与安全领域关系密切了

5、样本分析

恶意代码分析是整个威胁情报生产过程中的重要环节。绝大多数的文件分析是通过动静态的扫描分析引擎处理的，但某些新文件对象可能需要人工的介入，这也会是威胁分配中最耗费时间的环节。如果能利用包括大模型在内的 AI 技术帮助提升分析效率，可以取得非常好的效果。

下面是我们的一些尝试。

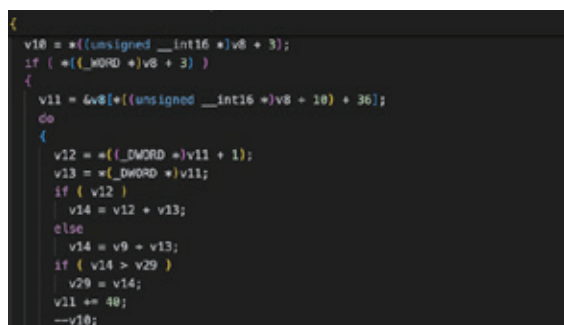
最早的大模型协助代码分析，以 IDA 插件的方式出现。它可以解析 IDA 中的函数，理解其功能给出可以描述性的函数名。这里是恶意程序分析二进制程序的反汇编代码中一个例子，另外，大模型分析脚本代码，哪怕是混淆后的，也非常得心应手。



WannaCry 勒索蠕虫样本变种
反编译 C 代码 247K
ed01ebfbc9eb5bbea545af4d01bf5f1071661840480439c6e5babe
8e080e41aa_diskpart.c

此前 Google 发布了一篇博客文章，介绍其 Gemini 大模型在分析恶意代码方面的亮眼表现。Gemini 1.5 Pro 提供了一次性 100 万 Token 的处理能力，分析反汇编和反编译的代码文本，一个中等规模的可执行文件处理时间仅为 30~40 秒，可以判断是否恶意并给出相应的参考信息，生成相当有用的报告。From Assistant to Analyst: The Power of Gemini 1.5 Pro for Malware Analysis (<https://cloud.google.com/blog/topics/threat-intelligence/gemini-for-malware-analysis>)

文章中提到两个样本例子，一个是 WannaCry 勒索蠕虫病毒的变种，另一个是在 VirusTotal 0 检出的恶意文件。我们找到那两个样本，用 IDA 生成了对应的反编译 C 代码。



VirusTotal 0 检出恶意样本
反编译 C 代码 357K
719b44d93ab39b4fe6113825349addfe5bd411b4d2508191656
1f9c403599e50_medui.c

通过发现可疑的加密的操作，判断样本偏恶意，内置的知识也提示与 WannaCry 有关。

第二个 VirusTotal 上的多病毒引擎的 0 检出样本对大模型是个高难度的考验，这里是 Kimi 和通义千问的表现。

在分析反编译的可执行代码时，我们需要深入理解代码的逻辑和行为。以下是对提供的代码的分析，包括其行为和潜在的恶意功能。

判定：

- 初步判定：恶意
- 原因：代码中存在多个可疑行为，如频繁调用系统函数、异常处理、数据加密等，这些通常是恶意软件用来隐藏其行为、绕过安全检测的手段。

活动列表：

IOC 列表：

- 无明显的外部IP地址或域名调用。代码可能主要在本地执行。
- 存在多个子程序调用，但没有明显的外部标识符。

创建的文件：

- 代码中没有直接显示创建文件的操作。但有一些函数可能与文件操作相关，如 sub_14000C450、sub_14000C620等。

注册表项：

- 代码中没有直接显示修改注册表的操作。

互斥体：

- 没有明显的互斥体创建或使用的代码。

网络活动：

- 没有直接的网络请求或套接字操作代码。

详细分析：

- 异常处理：**代码中存在异常处理的模式，如sub_1400010E0、sub_140001080等。这可能是为了捕获和处理执行过程中的异常，避免程序崩溃。
- 数据加密：**部分函数可能涉及数据加密或解密，如sub_140001510、sub_1400015F0等。这可能是为了隐藏恶意数据或行为。
- 系统调用：**代码频繁调用系统API，如OpenClipboard、GetClipboardData、SetClipboardData等。这表明程序可能与系统剪贴板交互，可能用于数据窃取。
- 进程和线程操作：**存在对进程和线程的操作，如K32EnumProcesses、OpenProcess等。这可能用于监视或操作其他进程。
- 隐藏行为：**代码中存在多个noreturn函数，如sub_140001220、sub_140001270等。这可能是为了隐藏其恶意行为，避免被用户或安全软件发现。
- 自删除机制：**部分函数可能涉及自删除机制，如sub_14000C590、sub_14000FDB0等。这可能是为了在执行后清除痕迹。
- 持久性机制：**没有明显的持久性机制。如注册表自启动项的创建。
- 反调试和反沙箱：**代码中没有明显的反调试或反沙箱技术。

结论：

根据上述分析，该代码表现出一些恶意软件的特征，如异常处理、数据加密、系统剪贴板交互等。然而，由于缺乏直接的网络活动、持久性机制和明显的外部通信，不能完全确定其恶意性。建议进一步分析和监控程序的行为，以确认其真实意图。

Kimi_2

2 份文档阅读完成

分析了提供的反编译代码后，这里是一份关于其行为的概述和分析摘要，以及是否具有恶意属性的判断：

判定

行为倾向： 良性

摘要： 代码主要涉及重复输出字符串“Microsoft Office assistant OnLine”，没有明显的破坏性或恶意行为特征。如网络通信、文件篡改、注册表操作等。然而，频繁的重复输出字符串可能在消耗系统资源，如CPU或内存空间，从而造成某种程度的性能服务效率。尽管如此，这种行为相较于传统意义上的恶意软件风险较低，但仍可视为一种潜在的资源浪费行为。

活动列表

- 输出大量信息：通过 sub_140000020 函数重复打印“Microsoft Office assistant OnLine”，共计多次。这种行为可能意在不断向用户界面、日志文件或控制台输出信息，造成干扰或占用资源。
- 无明显恶意活动：未发现直接修改文件、注册表操作、网络通信、互斥体创建等典型的恶意活动。

IOC 列表

- 无恶意IOC：由于没有网络通信、特定文件读写或注册表修改，没有直接的IOC可列出。

结论

虽然代码不表现典型恶意软件的严重特征，如数据窃取、远程控制或系统破坏，但其异常的输出行为可能消耗系统资源，特别是如果该设计为循环执行或在高频率下，可被视为轻微的骚扰或服务干扰。因此，虽然判断为良性与恶意之间存在灰色地带，但鉴于其潜在的负面影响，倾向于将其标记为可疑的恶意或至少是有问题的。进一步调查应关注代码如何被调用及其实际环境影响，以评估其造成的影响程度。

通义千问_2

总体来看，结果并不让人满意，Kimi 虽然判定为恶意，但并没有给出靠谱的理由，通义千问就直接判定来非恶意，不奇怪，毕竟这两个只是并没有根据安全分析需求做过调整的通用模型。之前对于第一个 WannaCry 变种样本的判定很可能更多依赖已经学习到老知识。

奇安信安全大模型的分析，也做出了准确的恶意性判定，并且给出了合理的依据。

该代码片段展示了恶意软件的特征，包括多种系统调用和文件操作，旨在分析其行为并判断其是否属于恶意软件。代码通过调用系统API实现文件操作、网络通信、进程管理、注册表操作、互斥体操作、网络活动、详细分析、结论。

1. 文件操作和系统调用

- 调用CreateFile函数，尝试打开或创建文件，如“C:\Program Files\Microsoft Office\Office\1033\excel.exe”。
- 调用WriteFile函数，向文件写入数据。
- 调用CloseHandle函数，关闭文件句柄。
- 调用GetLastError函数，获取最后发生的错误。

2. 网络通信

- 调用Socket函数，创建套接字。
- 调用connect函数，连接到远程主机。
- 调用send函数，向远程主机发送数据。
- 调用recv函数，从远程主机接收数据。
- 调用closesocket函数，关闭套接字。

3. 进程管理

- 调用GetCurrentProcessId函数，获取当前进程的ID。
- 调用GetProcessId函数，获取指定进程的ID。
- 调用OpenProcess函数，打开指定进程。
- 调用TerminateProcess函数，终止指定进程。

4. 注册表操作

- 调用RegOpenKeyEx函数，打开注册表项。
- 调用RegSetValueEx函数，设置注册表值。
- 调用RegCloseKey函数，关闭注册表项。

5. 互斥体操作

- 调用CreateMutex函数，创建互斥体。
- 调用WaitForSingleObject函数，等待互斥体。
- 调用ReleaseMutex函数，释放互斥体。

6. 网络活动

- 调用InternetConnect函数，连接到互联网。
- 调用HttpOpenRequest函数，创建HTTP请求。
- 调用HttpSendRequest函数，发送HTTP请求。
- 调用HttpReadResponse函数，读取HTTP响应。

7. 详细分析

- 代码中存在异常处理模式，如try-catch块。
- 代码中存在数据加密或解密操作，如xor函数。
- 代码中存在系统调用，如MessageBox函数。

8. 结论

- 根据上述分析，该代码表现出一些恶意软件的特征，如异常处理、数据加密、系统剪贴板交互等。
- 然而，由于缺乏直接的网络活动、持久性机制和明显的外部通信，不能完全确定其恶意性。
- 建议进一步分析和监控程序的行为，以确认其真实意图。

QAX-GPT 安全大模型_2

我们把 Google Gemini 模型的结果再拿回来做下对比：



Gemini Pro

剪贴板监控、查找进程、大量的字符串操作这些疑似恶意的行为两个模型都发现了，奇安信模型还发现了查找窗口这种明显恶意的操作，IOC 获取方面各有千秋。

小结

- 从分析能力上看，即便是没有专门为恶意代码分析做过优化的通用模型也表现出了一定的恶意代码描述能力，但还远远不够，经过安全分析能力调校的模型会有好得多的表现
- 上下文窗口的大小非常重要，直接影响一次性处理代码块的大小，影响分析的完整性，从而影响分析的准确度，但目前开放可用的大模型可能由于资源限制大多只能提供最大 128K 的上下文的窗口，需要通过一些技术方案来解决这个问题

6、威胁实体抽取

从非结构化的文档里抽取威胁元素实体及关系，这是威胁情报信息收集的关键环节。每天，我们从开源渠道能获取到大量的文章，描述各类新近的攻击活动，包含工具、过程、

手法及 IOC 信息，我们需要自动化的手段来高效完整地提取关注的内容。

以前，比较成熟的方法是基于神经网络的自然语言处理，模型的训练需要海量的数据，算法的开发应用需要相当专业的人员，使用还是比较麻烦，现在有了大模型要方便多了。

现在大多数的威胁活动分析文章会把 IOC 统一以相对结构化地呈现放在文章末尾。例子里的这个文章包含的 IOC 类型非常丰富。

Network				
IP	Domain	Hosting provider	First seen	Details
N/A	THEGUESTOWER.BY. master-gna- elcloud[.]net	N/A	2020-02-01	Domain (Free DNS) pinged by malicious Word macro.
45.33.24[.]145	N/A	Akamai Connected Cloud	2020-05-20	C&C server of LunarWeb (compromised VPS)
45.79.93[.]87	N/A	Akamai Connected Cloud	2020-05-20	C&C server of LunarWeb (compromised VPS)
65.109.179[.]147	N/A	HEIZNER Online GmbH	2023-10-29	C&C server of LunarWeb (compromised VPS)
74.90.80[.]35	N/A	Host Department NLLLC	2023-10-29	C&C server of LunarWeb.
82.145.158[.]186	N/A	IONOS SE	2022-08-03	C&C server of LunarWeb (compromised VPS)
82.223.111[.]220	N/A	IONOS SE	2022-08-03	C&C server of LunarWeb (compromised VPS)
139.162.23[.]113	N/A	Akamai Connected Cloud	2023-06-15	C&C server of LunarWeb (compromised VPS)
159.229.102[.]180	N/A	Contabo GmbH	2023-10-29	C&C server of LunarWeb.

IP 列表

Files			
SHA-1	Filename	Detection	Description
DE83C2C3FE49C818F96173E9EE36A5161DCFB24A	App_Web_0cm4b1br.dll	MSR/Agent.ERT	Compiled version of ASP.NET web page that installs LunarWeb.
9CEC3972FA35C80DE87BD6655DE18B3EDM6D877C	N/A	VBA/TrojanDownloader.Agent.ZJC	Malicious Word macro that installs LunarMail.
2ED792E39F72B6D852BD7F6AED96AFC93617897CF	qppol.dll	Win64/LunarLoader.B	LunarLoader (x64) used to load LunarMail.
FC9636553E486D042E80D9B25A390C0757FF8DE0	qppol.dll	Win32/LunarLoader.A	LunarLoader (x86) used to load LunarMail.
94A6C9C75B0C847878855896C41336770909414	tapiperf.dll	Win64/LunarLoader.C	LunarLoader (x64) used to load LunarWeb.
0006830836F915911349C28E81AE39325AD84	admpwd.dll	Win64/LunarLoader.A	LunarLoader (x64); a trojanized AdmPwd, used to load LunarWeb.
15D8CF2ED062BAE23ED1970EFA8BDAFC69552E85	AdmPwd.dll	Win64/LunarLoader.A	LunarLoader (x64); a trojanized AdmPwd, used to load LunarWeb.
755C4127D427F8D0FAF4510B46B52A588D8B03A	wlanet.dll.mui	Win64/LunarLoader.B	LunarLoader (x64) used to load LunarWeb.
7547B65715643FD03A680C9FC124528570CA80C	N/A	Win32/LunarWeb.A	LunarWeb backdoor (x86).

文件 Hash 列表

我们可以构造 Prompt 来让大模型来干活，Prompt 的设计是一个反复测试优化的过程，根据反馈的结果提要求加限制，最终得到我们想要的信息与格式。

假设你是一位网络安全专家，下面以 3 个单引号包含的网页链接是一个网络攻击活动的报告，请从中提取如下这些信息。

- 1) 发布这个报告的厂商。
- 2) 攻击者的名称，包括可能的别名。
- 3) 遭受攻击的目标。
- 4) 攻击目标的行业，需要尽可能详尽。

- 5) 攻击发生的起始和结束时间。
 - 6) 攻击者所使用的网络基础设施，包括：IP、域名、文件 Hash、URL、数字证书、注册表项。
 - 7) 攻击者所使用的软件或工具，需要尽可能详尽。
 - 8) 攻击者所使用的恶意代码家族，需要尽可能详尽。
 - 9) 攻击者所使用的硬件。
 - 10) 攻击者所使用的攻击手法，添加手法对应的 ATT&CK 技术点 ID。
 - 11) 如果有相关的信息的话，生成相应的流量检测规则
 - 和文件检测规则。
 - 12) 攻击者所使用的安全漏洞，包括 ID 或漏洞名。
- 输出中不能包含文本中不存在的对象。
- 以 JSON 对象的形式给出结果，全部使用英文字符集，不存在或未知的值设置为空，以方便拷贝的代码格式输出。

“ ”

https://www.welivesecurity.com/en/eset-research/moon-backdoors-lunar-landing-diplomatic-missions/

“ ”

我们调用了 Llama3 模型执行这个抽取任务，结果看起来很不错，在这个案例里 TTP 和 IOC 信息抽取得很完整，而且自动去除了混淆（如 IP 字符串中的方括号）。

输出的信息需要保持字段和结构的一致，这可以通过给出样例的方式来约定，聪明一些的大模型会严格按给定的样例来输出。



Prompt 在原来的基础上增加了输出样式的指定，Llama3 这样级别的模型完全可以正确处理。

小结

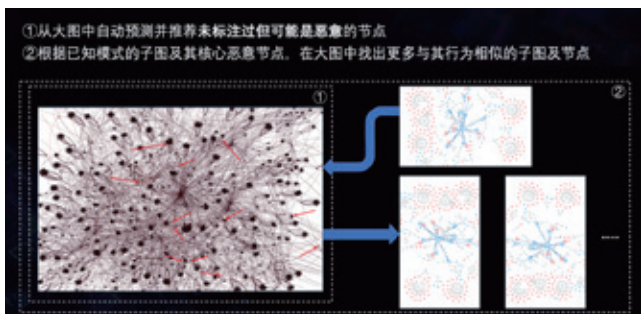
- 输出样式如果没有明确的示例会非常随机
- 数据提取不能保证完整，因为有 Token 窗口限制，并且不一定每次保持一致
- 虽然有明确要求，但还是可能出现幻觉，后续如要严肃使用需要加入校验环节
- 模型平台具备处理网页链接的能力会非常方便，而且可能节约 Token 量
- 对于链接的爬取不是非常稳定，可能由于网络原因爬取失败

7、大规模威胁知识图谱的构建和应用

威胁情报的生产涉及海量的多类型数据的分析，典型的有 IP、域名、文件 Hash、URL、数字证书、邮箱地址等，这些实体之间从安全分析的角度可以被组织成知识图谱。构建完成以后可以作为后台被用来支持各种应用，如自动化的关联性检索，交互式的实体探索，基于机器学习的图相似分析。

目前，在这个方向上我们还没使用到大模型，但在学术圈已经有了些尝试，以后有可能应用到工业界。

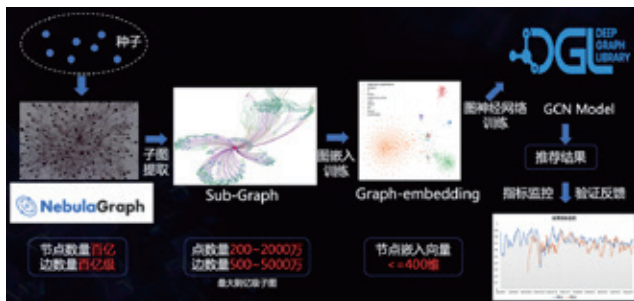
基于已知恶意节点，在我们的数据视野范围内拓展发现在更多的恶意节点，这是我们威胁情报运营中的一个非常核心的需求。利用威胁实体知识图谱，结合机器学习实现可疑的恶意节点推荐是目前我们一个比较成功的实践。



为此，构建了一个目前规模已经达到百亿节点与边的图数据库，节点类型超过 20 种，经过子图提取和图嵌入训练，图神经网络训练，通过模型输出推荐的疑似恶意节点。

我们构建了一个目前规模已经达到百亿节点与边的图数

据库，节点类型超过 20 种，经过子图提取和图嵌入训练，图神经网络训练，通过模型输出推荐的疑似恶意节点。



基于图谱的推荐，已经形成稳定的情报运营流程，每日控制推荐 100 个未知黑节点提交人工确认，准确率超 50%。看起来准确率不高？运营的核心要点在于根据可投入的人工验证资源多少控制每日绝对输出量，只要绝对数量可以做到日清，50% 的准确率是完全可接受的。



一个推荐系统输出的可疑，最终确认恶意的例子

8、智能 Agent

以上各环节中采用的 AI 及大模型技术主要用于一些单个的点。接下来我们说一下，当前最热的基于大模型的智能体技术。

要完成一个相对复杂的任务，需要把各个流程串联起来，目前已经出现大量的 Agent，在安全方面的应用只是其中一个小角而已。下面这个问答应用就是奇安信内部一个比较大的 Agent 项目，界面上看起来简单，但后台的流程其实相当复杂，从意图识别，安全检测过滤，到后台各类数据源的对接，要实现好的反馈效果与性能，需要大量的工程化设计与测试调试。



<https://qgpt.qianxin.com/chat>

目前，我们还在探索其他威胁分析类的场景。在 Agent 方向，我们可以做的事情还很多，如情报的更自动化的关联分析。



9、报告生成

威胁情报生产出来以后，除了将数据提供给检测类设备消费，还有给安全人员看的。

在这方面，模型的文本总结和摘要能力可以被充分利用起来。奇安信 Alpha 平台上，对一个域名的研判页面，罗列了各类维度的信息。



其实这样的信息陈列对一般的用户来说结论还不够直接。生成一个人读的总结报告是更友好的界面，怎么做最容易？

最简单的方案，整合一个数据接口出来，包含一系列结构化的核心信息，前面定制一个 Prompt，交给大模型整合生成总结就行了，本质上一个最简化的 RAG，结果远好于一般的安全分析人员人工撰写。

Prompt 加部分数据的例子：

“假设你是一个网络安全专家，以下这个 json 对象是关于一个网站域名（data.whois_abstract.domain）相关的情况。首先，请输出域名相关的基本信息，包括：注册信息（data.domain_attribute.createdate,data.domain_attribute.updatedate,expiredate），是否为动态域名（data.domain_attribute.is_dynamic），是否为白名单域名（data.domain_attribute.in_whitelist），当前绑定过的域名（data.subdomain.examples），当前的标签信息（data.domain_tag.examples）。

然后，请基于如下这些维度信息对域名是否倾向于恶意做判定：

- 标签信息（data.domain_tag.examples）
- 域名属性（data.domain_attribute.malicious）
- 域名相关的文件样本（data.domain_try_connect.examples）

是否恶意的结论写在输出的前部，并给出相应的理由。如果存在域名相关的文件样本（data.domain_try_connect.examples），请罗列出来详细的信息。

最后，给出相关的参考资料链接（data.report.examples），展示具体链接文本，并下载解析链接对应的

内容并做总结。

不要提 json 里的具体字段名。请全部用中文答复。

```
{
  "status": 10000,
  "message": " 执行成功",
  "data": {
    "domain_attribute": {
      "popular": "",
      "createdate": "2022-06-07 15:10:40",
      "updatedate": "2023-08-19 19:03:10",
      "expiredate": "2023-06-07 15:10:40",
      "last_seen": "2024-03-21 19:07:20",
      "is_dynamic": false,
      ....
    }
  }
}
```

我们来看看各模型对于一个域名进行总结的输出表现。
对它们的输入，从 Prompt 到数据都是完全一样的。

域名基本信息

- 注册日期: 2022-06-07 15:10:40
- 更新日期: 2023-08-19 19:03:10
- 过期日期: 2023-06-07 15:10:40
- 是否为动态域名: 否
- 是否为白名单域名: 否
- 当前绑定过的域名: www.cilis-cn.net, m.cilis-cn.net, ftp.cilis-cn.net
- 当前标签信息: APT, APT-S-0020, KNOWN APT, C&C, APT-C-09, Dropping Elephant, PatchWork, 摩可草, 白象, Patchwork_16_07_08

是否倾向于恶意的判定

结论: 该域名倾向于恶意。

理由如下:

- 标签信息中包含多个与APT (高级持续性威胁) 相关的标签, 如"APT", "APT-S-0020", "KNOWN APT", "APT-C-09", "Dropping Elephant", "PatchWork", "摩可草", "白象", "Patchwork_16_07_08"。
- 域名属性中标记为"AP (恶意)"。
- 域名相关的文件样本中存在恶意类型为"Trojan"的样本, MD5 为"953fee8ef679f8c6b6a19f6fdb0ae9a1"。

域名相关的文件样本详细信息

- 恶意类型: TROJAN
- MD5哈希: 953fee8ef679f8c6b6a19f6fdb0ae9a1
- 时间戳: 2017/04/14 03:30:31
- 家族: selfdel

参考资料链接

1. [LevelBlue - Open Threat Exchange](#)

该链接指向AllerVault的Open Threat Exchange, 可能包含有关该域名的威胁情报和分析。

参考资料内容总结

- 由于提供的参考资料链接内容未在JSON中给出, 无法进行下载和总结。这些资源会提供关于域名的更多背景信息, 包括安全研究人员的分析, 历史行为和其他可能的威胁指标。

综上所述, 基于域名的标签信息、属性以及相关的文件样本, 可以判定该域名具有恶意倾向, 并多个已知的APT组织和恶意活动相关联。建议采取必要的安全措施以防范可能的安全风险。

95

28

域名倾向于恶意。

基本信息

- 注册日期: 2022-06-07 15:10:40
- 更新日期: 2023-08-19 19:03:10
- 过期日期: 2023-06-07 15:10:40
- 是否为动态域名: 否
- 是否为白名单域名: 否
- 当前绑定过的域名: www.cilis-cn.net, m.cilis-cn.net, ftp.cilis-cn.net
- 当前标签信息: APT, APT-S-0020, KNOWN APT, C&C, APT-C-09, Dropping Elephant, PatchWork, 摩可草, 白象, Patchwork_16_07_08

判定理由

该域名包含APT (高级持续性威胁)、C&C (命令与控制服务器) 以及与知名APT活动相关的标签 (如PatchWork和Dropping Elephant)。这些标签是恶意活动的常见标识。

域名属性中的"恶意"标记为"AP (恶意)", 这是根据域名已知的APT活动相关的恶意行为。

- 存在域名相关的恶意文件样本, 具体信息包括:
 - 恶意类型: TROJAN (木马)
 - MD5哈希: 953fee8ef679f8c6b6a19f6fdb0ae9a1
 - 时间戳: 2017/04/14 03:30:31
 - 家族: selfdel (自删除木马)

相关的参考资料链接

- [Patchwork Analysis - Google Docs](#)
- [Patchwork Analysis \(TXT\) - Google Docs](#)
- [Alienvault OTX - Patchwork](#)

综合以上信息, 该域名与多个已知的APT活动关联, 并且存在与域名的已知恶意文件样本。因此, 该域名被判定为具有高风险。

GPT-4

域名风险评估结果: 高风险

基本信息

- 域名注册信息: 创建日期2022-06-07 15:10:40, 更新日期2023-08-19 19:03:10, 到期日期2023-06-07 15:10:40
- 是否为动态域名: 否
- 是否为白名单域名: 否
- 当前绑定过的域名: www.cilis-cn.net, m.cilis-cn.net, ftp.cilis-cn.net等
- 当前标签信息: APT, APT-S-0020, KNOWN APT, C&C, APT-C-09, Dropping Elephant, PatchWork, 摩可草, 白象, Patchwork_16_07_08等

风险评估

基于标签信息、域名属性和域名相关的文件样本, 我们认为该域名存在高风险的恶意行为。

标签信息中包含多个恶意标签, 如APT、C&C、PatchWork等, 表明该域名可能涉及恶意活动。

- 域名属性中标记为恶意 (malicious), 并且safe_level为50, 表明该域名存在一定的恶意风险。
- 域名相关的文件样本中包含恶意文件类型TROJAN, 表明该域名可能用于恶意活动。

域名相关的文件样本

- 恶意文件类型: TROJAN
- MD5: 953fee8ef679f8c6b6a19f6fdb0ae9a1
- 时间戳: 2017/04/14 03:30:31
- 家族: selfdel

参考资料链接

- https://docs.google.com/document/d/1oYX3uN6KxIX_StzTH0s0yFNNcHDrV8VgmVqU5WoeErc/mobilebasic
- https://docs.google.com/document/export?format=txt&id=1oYX3uN6KxIX_StzTH0s0yFNNcHDrV8VgmVqU5WoeErc
- <https://otx.alienvault.com/pulse/5a143066e6ecbc245110c200/>

下载解析链接对应的内容, 总结为: 该域名可能涉及恶意活动, 如APT攻击、命令控制服务器等。存在高风险的恶意行为。

Llama3 70B

结果看起来差不多? 其实, 从一些细节可以显示模型在安全分析方面的一些能力区别。在判定理由部分, Kimi 只是简单罗列了发现的标签; Llama3 则能把真正影响判定结果的标签给挑出来, 显然 GPT4 更强。GPT4 还能对标签做

出分类，理解 APT C&C 是什么意思，哪些标签是 APT 组织名，这是模型内置知识与推理能力的综合体现。

另一个文件样本的沙箱日志的例子。在行为活动的结构化数据前添加一个简单的 Prompt，交给大模型总结，我们可以对其自动生成人读报告。

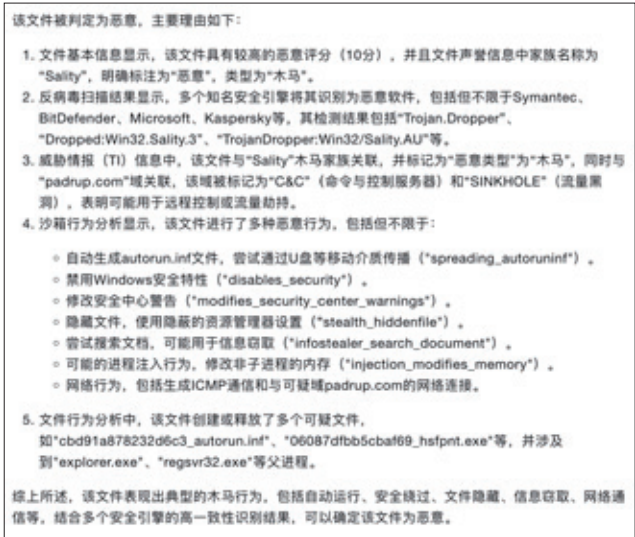


一个样本的沙箱行为展示



对应的 Prompt 及数据

与域名分析报告类似，给定部分样本的关键信息以后，大模型自动生成人读的报告质量基本可用。



对于漏洞信息方面，我们也可以更方便地基于多源信息进行总结，整合搜索最新信息与 AI 总结能力的新搜索

引擎，对传统漏洞库是降维打击般的存在。国外的应用有 Perplexity，国内的则有秘塔，还有一些开源项目也实现了类似的功能，甚至可以本地部署，缺点是调用搜索 API 成本并不低，也不能大规模地对外提供服务。

除了总结一些技术层面的描述，还可以通过搜索整理相关的安全事件，提示漏洞在现实世界里所导致的实际风险，对用户采取应对措施的优先级排序非常有指导意义。

CVE-2024-3400

漏洞

资讯

漏洞

205

漏洞

资讯

漏洞

CVE-2024-3400 是一个严重的零日漏洞，影响 Palo Alto Networks PAN-OS 软件中的 GlobalProtect 功能。该漏洞允许未经身份验证的远程攻击者在受影响的防火墙上执行任意命令。从攻击者角度看，该漏洞可被用于多种用途。

根据公开资料，这个漏洞被分配了 CVE-2024-3400 的编号，表明其危害性很高。此外，已有黑客组织如 UTA0210 积极利用该漏洞进行攻击。攻击者利用该漏洞“午夜行动”。

漏洞的技术细节包括通过漏洞利用程序和合法 ID 处理中的命令注入来允许任意命令写入。这可以通过通过漏洞利用程序的 Cookie 参数来触发。受影响的 PAN-OS 版本包括 10.2、11.0 和 11.1 等版本。

Palo Alto Networks 已于 2024 年 4 月 12 日发布了关于此漏洞的安全通告，并建议用户采取相应的缓解措施。他们提供了缓解措施和补丁信息。此外，公司还提供了关于如何缓解此漏洞的指导，包括使用专用数据流和配置防火墙策略。

总结来说，CVE-2024-3400 是一个高危的命令注入漏洞，影响特定版本的 PAN-OS 软件中的 GlobalProtect 功能。用户应尽快采取缓解措施，以防止潜在的安全风险。

漏洞描述

漏洞影响

漏洞利用

漏洞修复

漏洞缓解

漏洞利用

漏洞修复

漏洞缓解

相关事件

相关人物

事件名称	事件时间	事件描述
CVE-2024-3400 零日漏洞发现并公开(安全)	2024年4月12日至2024年4月29日	根据公开资料，CVE-2024-3400 漏洞是在 2024 年 4 月 12 日被发现的。该漏洞允许未经身份验证的远程攻击者在受影响的防火墙上执行任意命令。
CVE-2024-3400 漏洞影响范围和严重性(安全)	2024年4月12日至2024年4月29日	根据公开资料，CVE-2024-3400 漏洞影响 Palo Alto Networks PAN-OS 的特定版本。该漏洞被分配了 CVE-2024-3400 的编号，表明其危害性很高。
Palo Alto Networks 发布了安全通告(安全)	2024年4月12日至2024年4月29日	Palo Alto Networks 已于 2024 年 4 月 12 日发布了关于此漏洞的安全通告，并建议用户采取相应的缓解措施。

相关人物

漏洞

秘塔
<https://metaso.cn/>

CVE-2024-3400漏洞相关的信息

来源

漏洞

漏洞

漏洞

2

漏洞

漏洞

漏洞

CVE-2024-3400 漏洞描述

CVE-2024-3400 是一个严重的命令注入漏洞，影响 Palo Alto Networks 的 PAN-OS 软件中的 GlobalProtect 功能。该漏洞允许未经身份验证的远程攻击者在受影响的防火墙上执行任意命令。

漏洞影响范围

该漏洞影响以下版本的 PAN-OS，并且需要启用 GlobalProtect 网关或门户以及设备端功能：

- PAN-OS 10.2 < 10.2.9-M
- PAN-OS 11.0 < 11.0.4-H1
- PAN-OS 11.1 < 11.1.2-A3

其他版本的 PAN-OS 以及 Prisma Access、Panorama 设备或 Cloud NGFW 不受此漏洞影响。

漏洞利用细节

攻击者可以利用该漏洞建立每分钟执行一次的 cron 作业，从外部服务器获取命令并执行。研究人员发现攻击者上传了一种名为 UPSTYLE 的 Python 后门程序，用于执行进一步的恶意活动。

修复建议

官方已发布安全更新补丁，建议受影响用户尽快升级到安全版本。同时，可利用安全功能如可信地址和防火墙规则。

关于此漏洞有哪些安全事件?

来源

漏洞

漏洞

漏洞

2

漏洞

漏洞

漏洞

午夜行动概述

自 2024 年 3 月 26 日起，一个身份不明的黑客组织一直在利用这个 Palo Alto Networks PAN-OS 软件中的零日漏洞 CVE-2024-3400。

攻击者利用该漏洞建立每分钟执行一次的 cron 作业，从外部服务器获取命令并执行 bash shell 运行。

研究人员发现攻击者上传了一种名为 UPSTYLE 的 Python 后门程序，用于执行进一步的恶意活动。

攻击者还利用了与防火墙相关的漏洞文件 (/var/log/pan/sshvpn_err.log 和 /var/log/pan/sshvpn_err.log) 来记录攻击行为。

攻击者的目标似乎是受影响的设备中获取敏感信息。

综上所述，午夜行动是一起由利用 CVE-2024-3400 漏洞的未知威胁行为者发起的大规模网络攻击活动，旨在窃取 Palo Alto Networks 防火墙并执行恶意行为。

关于此漏洞有哪些安全事件?

来源

漏洞

漏洞

漏洞

2

漏洞

漏洞

漏洞

Perplexity
<https://www.perplexity.ai/>

其实，使用 Kimi 处理漏洞相关的参考信息链接也能达到类似效果。当然，漏洞相关的结构化关键信息还是需要的，

通过 API 集成到组织内部的漏洞运营流程里。



简单的 Prompt 加上相关的网页链接



Kimi 会爬取链接的内容并做总结输出

我们已将上述这些对威胁实体的大模型总结集成到奇安信 Alpha 威胁研判分析平台，这对用户的使用体验是个很大的提升。从现在开始，用户只要登录访问奇安信威胁分析研判平台，直接就能看到集成的 AI 总结。



10、RAG 威胁信息检索增强

最后，说一下目前也非常热门的 RAG，这是结合大模型的生成能力与本地或新近内容提供诸如问答服务的系统，通过外挂知识的方式，解决私有数据和实时新信息的检索，这个能力在网络安全领域也是真实的需求。

事实上，大模型在安全检测和分析系统中的最早应用，如 SOC、NDR 等系统，其核心就是本地的 RAG 系统，为用户提供自然语言的问答接口，大模型分析用户需求的任务类型，基于模型自身的知识或转化为内部各类数据源的查询获取相关信息，最后，大模型对于答复的信息生成人读的答复，这也是加强安全设备使用检验的最基本应用。

在威胁情报运营输出方向，我们也做了些实践。这里的例子是尝试把一个国外的威胁百科网站转化为本地可问答查阅的系统。

Malpedia 是一个网络威胁信息的百科，包含 15,000 多个文章的链接，提供链接下载。可以尝试将其指向的内容 RAG 化，支持基于自然语言的搜索和关联。



https://malpedia.caad.fkie.fraunhofer.de/library

我们尝试了 QAnything 工具，把从链接指向的内容通过 jina.ai 平台爬取下来导入本地知识库，但可能限于机器的性能，效果并不太理想。

不过，目前类似的 RAG 工具越来越多，也越来越成熟。如果我们采用更好的硬件，配合更完善的软件，应该可以得到非常好的结果，这也是迟早的事。

三、实践总结

1. 大模型作为通用智力引擎几乎可以赋能所有的任务，组织需要在内部和外部构建这样的智力基础设施。本地部署需要一定的硬件资源投入，主要应对处理私有数据的场景，效费比未必高；如果只是处理公开非敏感数据，建议直接使用各大模型厂商的云端服务，效果可以，费用其实也不高。
2. 大模型技术可以极大提升安全分析及威胁情报运营的效率，降低运营流程开发人员的能力要求，以较小的人工投入获得更好的结果，在各环节中综合使用 AI 技术可以提升 50% 的效率。
3. 长上下文窗口能力很关键，对于文本内容的总结、恶意代码的分析、漏洞的挖掘、RAG 功能有决定性的影响，需要持续开发相关的技术。
4. 大模型自身能力的强弱会极大影响提示词工程和产品化的工作量，背后的道理很简单，让聪明人干活你只要给出

目标，不需要指导怎么做，但对能力不强的下属，你甚至得指导具体的执行步骤。

5. 各个大模型在不同的任务中各有擅长的表现，目前还没发现全能的模型同时满足我们的所有需求：

- Kimi 的中文总结与指令跟随能力强，上下文 Token 空间大，能直接处理链接爬取，通义千问综合能力差不多
- GPT4 的知识和推理能力最强，中文处理尚可
- LLaMA3 推理能力接近 GPT4，中文处理能力差（原始模型），但最大的优势作为开源模型可本地部署
- 特定的安全分析任务，如恶意代码分析、攻击检测，通用模型的能力还是不够，需要专用的安全分析模型

6. 基于大模型技术的各类包装应用，RAG、LangChain、DIFY、Agent 等风起云涌，很多时候并不需要自己开发，直接使用改进即可

7. 本地 RAG 系统涉及很多工程细节及原始数据的质量，要出好的效果比较有挑战

8. 警惕幻觉，哪怕是 RAG 和基于限定内容的总结，也可能出现，基于 GenAI 的判定对外执行操作需非常谨慎 安

关于作者

汪列军

虎符智库专家，奇安信威胁情报中心负责人。

奇安信威胁情报中心

中国威胁情报行业领军者

奇安信威胁情报中心是奇安信集团旗下专注于威胁情报收集、分析、生产的专业部门，以业界领先的安全大数据资源为基础，基于奇安信长期积累的威胁检测和大数据技术，依托亚太地区顶级的安全分析师团队，通过创新性的运营分析流程，开发威胁情报相关的产品和服务，输出威胁安全管理与防护所需的情报数据，协助客户发现、分析、处置高级威胁活动事件。

威胁研判分析平台ALPHA： 一站式云端SaaS服务的威胁分析工具平台。是安全分析师为同行打造的利器，针对IOC查询、线索关联、事件溯源、样本行为检测和同源分析等威胁分析场景提供全面的解决方案。

高对抗云沙箱分析平台： 提供多种动静态检测、分析技术，展现文件各方面特征，帮助分析师快速判别、并掌握恶意软件的详情。

威胁分析武器库： 服务于安服、安运、安全分析师及各类企业用户。支持IOC自动化数据流检测、失陷情报、恶意IP批量查询；支持邮件批量自动化检测；支持WEB应急日志、主机应急日志分析。

威胁情报系统QAX TIP： 威胁情报平台是情报收集、维护、使用的综合性平台。可帮助企业在安全运营中，利用威胁情报快速检测、响应、分析和预防各类网络攻击威胁，并分析产生行业威胁情报。

威胁雷达： 利用大数据和威胁情报监测技术，整合了奇安信的高、中位威胁情报能力，提供指定区域内网络威胁高位监控、案件拓线分析、样本深度分析等多种能力。

样本同源分析系统： 奇安信威胁情报中心红雨滴团队基于样本基因深度解析，使用机器学习算法用于APT数字武器追踪的可视化分析系统。

高级威胁分析服务： 为网络安全主管单位、政府机构、行业客户及高校科研机构提供情报线索拓线和威胁分析服务，输出深度分析报告供其决策参考。

奇安信威胁情报中心：
ALPHA网址：<https://ti.qianxin.com>
雷达网址：<https://r.ti.qianxin.com>
扫描关注我们的微信公众号
邮箱：ti_support@qianxin.com



从零建设，看通信服务商如何打造全方位的网络安全“免疫力”

作者 营销中心服务部 鄧学兵

2022年5月，通信服务商G公司成立，致力于提高全国网络速度，成为推动科技进步的领军企业，经营范围遍布31个省。

不同于传统企业可以依托已有网络基础设施进行升级迭代，G公司必须从零开始规划与建设一个既能满足当前需求，又能适应未来发展的核心网络体系。在当前数字化转型加速、网络攻击手段日益复杂多变的背景下，网络安全的重要性不言而喻，尤其是G公司的门户网站、充值系统等，直接面向客户的关键业务平台，一旦发生数据泄露或被恶意攻击，不仅会损

害客户利益，甚至造成其他严重影响。

从零建设，规模大、时间紧、任务重

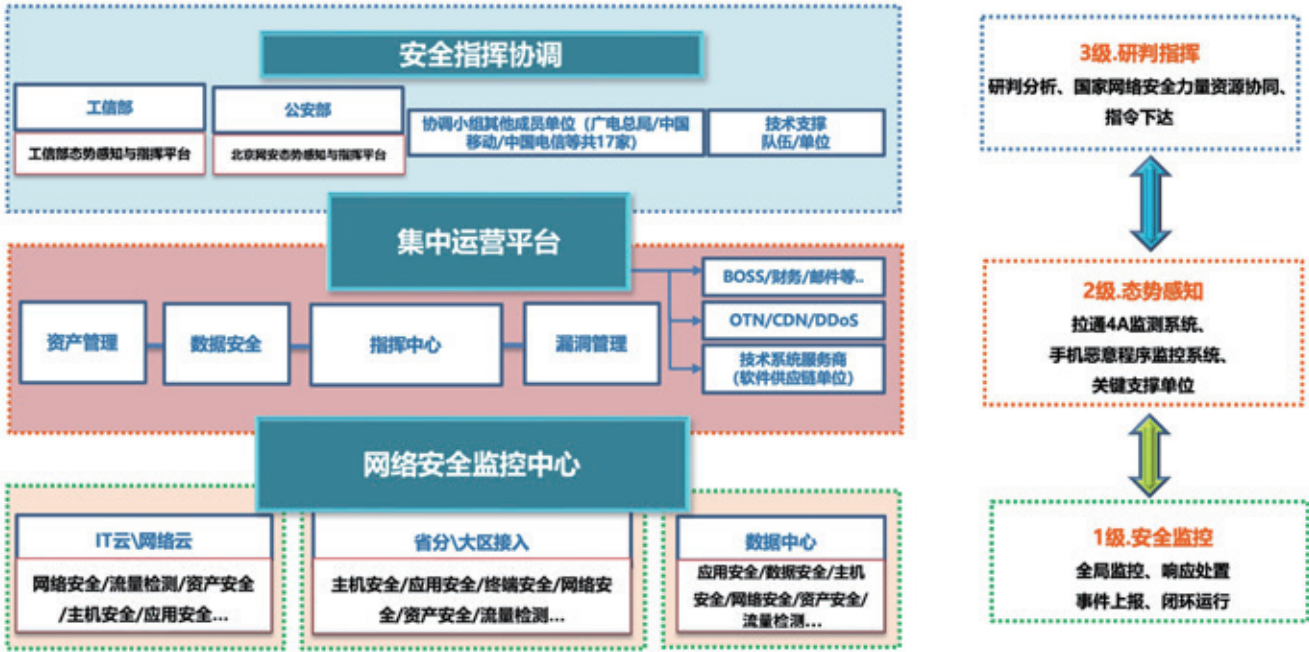
G公司安全总监表示：“公司高度重视网络安全的建设工作，将其视为核心竞争力的一部分。然而，作为新成立的企业，我们在网络安全方面的运营人员相对有限，特别是在需要跨31个省份协调资源、统一部署安全策略时，所面临的难度尤为突出。因此，我们迫切需要一家经验丰富、有过大型项目成功案例的网络安全服务商作为合作伙伴，希望能在半年内帮助我们构建一套既能抵御外部恶意攻击，又能高效监测内部风险的网络安全体系。”

鉴于奇安信冬奥“零事故”成功经验，以及在大型企业服务中所展现的专业性和可靠性，G公司委托奇安信建设一套网络安全集中运营分析平台，确保网络日常运营安全稳定，并在工信部考核及重大突发事件时，能够迅速实现信息的互联互通与高效指挥调度，构筑起一道坚不可摧的安全防线，为G公司基础设施提供全方位、综合性的保护。

项目启动之初，奇安信集团高度重视、迅速行动，从全国31个省份中



集团统一指挥协调，集中安全运营，三级平台联动



抽调出一批技术精湛、经验丰富的专家，组成项目领导团队与实施团队。在深入了解具体需求后，奇安信团队立即开展全面的需求调研，对项目的关键工期、质量标准、功能需求及可能遇到的挑战和问题有了清晰的认识。基于此，团队进行了细致的工作分解，制定详尽的 SOW、进度计划和实施方案等，确保每个环节有序推进。同时，团队还针对当时新冠疫情对项目进度带来的风险进行深入讨论，制定了预案措施，以确保项目能够按时、按质完成。

构建系统化的内生安全能力

经过详细调研，奇安信团队确认 G 公司想要打造完备的网络安全运营服务体系，以及对公司资产进行有效

的管理和安全防护。经过 6 个月的奋战，奇安信完成网络安全态势感知平台建设，帮助 G 公司构建网络态势感知能力和指挥调度体系，满足集中网络安全保障与服务运营需求，实现对网络基础设施的安全防护。通过搭建资产安全管理平台，完成对公司全网资产管理、资产排查、及时发现并修复资产漏洞，实现对资产的管理及安全防护。

网络安全集中运营分析平台建设分为总部节点和 31 省份节点。总部节点统一汇聚省级资产及态势感知基础数据，利用总部平台能力，打造纵向可与部委系统、省级平台信息共享，横向可与其他网络安全系统协同联动的一体化网络条线态势感知治理体系；省级平台重点实现资产信息采集、数据分析、数据展示、数据存储能力，满足工信部考核的同时，具备对省份网络安全的可视、可管、可控。

总部网络安全集中运营平台重点打造以下关键能力：

- 全面资产管理能力，能够全面地识别、分类、跟踪和管理组织内的所有网络资产，包括硬件、软件、网络设备等；
- 安全态势感知能力，通过收集多元、异构的海量数据，利用关联分析机器学习、威胁情报等技术，提供全局风险评估和应急响应的决策支持，协助安全运营人员进行威胁发现、调查分析及相应处置；
- 数据统一治理能力，实现数据汇聚与关联分析、统一化的基础数据处理、数据存储等，确保数据的一致性、可用性和安全性；
- 内生安全能力，网络安全集中运营平台设计了复杂异构环境下网络安全协同联动机制，形成了全生命周期网络安全部署体系，对系统内部安全

能力统一规划、分步实施，逐步建成面向数字化时代的一体化安全体系。

省级网络安全集中运营平台，实现一种或多种数据源的采集，并支持对数据源的增加、删除和修改管理。其中，数据源包括但不限于资产数据、日志数据、网络状态数据、归解析服务器的解析数据、网络攻击事件、恶意程序事件、威胁情报数据和公共漏洞发布平台的漏洞信息；具备数据处理、数据存储功能；具备安全态势预警及用户行为分析的能力。

实现全方位的网络安全“免疫力”

网络安全集中运营分析平台为 G 公司的核心网络建设了一套以“人 + 流程 + 数据 + 平台”为中心的安全运行体系，通过配置专业的安全运营人员、制定标准工作流程、规范各方协同机制，依靠平台开展日常的基础运

维、日志接入、告警监测、分析研判、响应处置、威胁建模、资产管理等工作，有效提升了 G 公司核心网络的网络安全防护能力，降低了潜在的安全风险。同时，平台的高效运行也进一步优化了 G 公司网络安全的运营架构，提高了事件及数据处理的速度和效率。

“奇安信从顶层设计的角度规划整体建设方案，系统性、全局性地进行核心网络安全的建设工作，从网络、数据、应用、行为、身份5个层面实现全网安全监测无死角、告警实时精准分析、安全事件高效处置，极大降低网络攻击风险，真正保证业务安全。”G公司安全总监表示。

共同迎接数字经济新时代

网络安全集中运营分析平台实现了对 G 公司核心网络安全保护的“实战化、体系化、常态化”，有效落实“主动防御、动态防御、精准防护、联防

联控”各项措施，真正做到对网络安全事件“事事有跟踪，跟踪有落实，落实有反馈”的闭环管理。

该项目对奇安信而言同样具有深远的战略意义。从方案设计、实施到长期运维的全方位服务，成为奇安信在拓展全国性大规模网络运营企业市场的里程碑事件。项目的成功实施，填补了奇安信在这一领域安全能力建设和全方位服务的空白，进一步巩固了奇安信在网络安全行业的领先地位。

随着 5G、物联网等新技术的快速发展，网络安全的需求将更加迫切。今年年初工信部印发通知，进一步部署了 2024 年信息通信业安全生产和网络运行安全工作，将完善制度政策体系、增强安全预防能力、提升应急处置水平、严格执法监督考核等七项列为重点任务。奇安信也将继续加大投入，不断创新，为通信行业提供更加全面和高质量的网络安全服务，共同迎接数字经济的新时代。安



Gartner：软件供应链安全指南

软件供应链安全是一个关键的风险和合规性问题，但大多数组织都以分散的方式处理它。缺乏一个包罗万象的框架会遗留安全漏洞。

主要发现

- 对软件供应链的攻击给组织带来重大的安全、监管和运营风险。有数据显示，这些攻击造成的损失将从 2023 年的 460 亿美元上升到 2031 年的 1380 亿美元。

- 在全球范围内，包括法律法规在内的合规要求，以及非正式的行业指导正在实施，迫使对软件供应链安全 (SSCS) 和应用程序安全风险采取更积极的应对措施。

- Gartner 2023 年技术采用调查发现，近三分之二的组织报告称他们已经实施或正在实施 SSCS 计划。尽管如此，多起事件和指标表明，这些努力（通常在整个组织内缺乏协调）未能解决严重的安全漏洞。

建议

- 制定一项协调战略来增强 SSCS，解决软件供应链框架的管理、创建和使用支柱中的所有风险因素。

- 确定安全、软件工程、采购、供应商风险管理和其他方的利益相关者，教育他们了解所涉及的风险并支持他们采取必要的行动以减轻危险。

- 与相关利益相关者协调，制定并

实施计划，以弥补与 SSCS 工作相关的流程和工具的不足。

战略规划假设

到 2027 年，80% 的组织将在整个企业内采用专门的流程和工具来降低软件供应链安全风险，而 2023 年这一比例约为 50%。

介绍

破坏性攻击凸显了对软件供应链安全的迫切关注。这些攻击的成本估计高达数百亿美元，预计到 2031 年将

增长 200%，达到 1380 亿美元（见图 1）。

主要攻击面和风险包括如下。

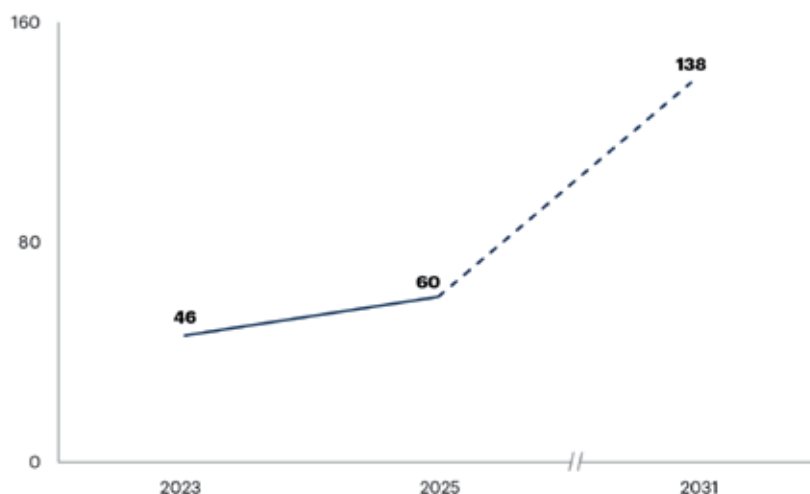
- 软件开发基础设施：通过渗透开发环境，攻击者将恶意软件整合到已部署的软件中。

- 恶意代码：2023 年在 245,000 个开源软件包中发现了恶意软件，是 2019 年以来所有年份发现总数的两倍。

- 开源漏洞：超过 90% 的专有代码包含开源依赖项，其中 74% 包含高风险依赖项。安全开发实践不足会增加代码中出现漏洞的概率。

- 被遗弃且维护不力的开源依赖

Global Annual Cost of Software Supply Chain Attacks
\$US Billions



Source: Cybersecurity Ventures/Snyk, October 2023
807046_C

Gartner

图 1：全球软件供应链攻击年度成本

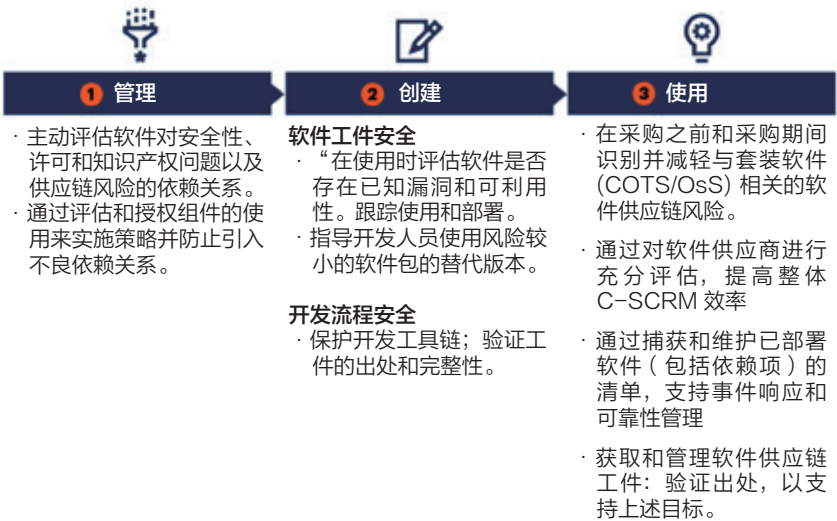


图 2 软件供应链安全三大支柱框架

项：大约一半的开源依赖项在两年多的时间里没有开发活动，导致这些项目实际上被遗弃了。许多项目缺乏强大的维护者团队。如果在这些依赖项中发现漏洞，修复要么会被延迟，要么根本就不会出现。

软件供应链安全也是一个监管和合规问题，全球政府和行业都提出了要求。

Gartner 技术采用报告的约 59% 受访者表示，他们已经或正在部署 SSCS 措施。但 Gartner 研究表明，这些努力往往脱节。努力往往集中在问题的单一方面，组织利益相关者之间几乎没有协调。这种孤立的

方法会导致重复和低效，以及糟糕的结果。许多组织正在重新考虑他们的软件供应链战略。在 Gartner Peer Connections 讨论中，67% 的参与者表示他们已经重新评估或计划重新评估他们的 SSCS 战略（截至 2024 年 5 月）。

软件供应链安全可以看作是一个涵盖三大支柱的框架：管理、创建和使用，如图 2 所示。通过实施这样的框架及支持流程和工具，安全和风险管理领导者可以确保对问题做出协调响应，最大限度地减少保护盲点或漏洞，并降低整个软件开发和使用生命周期的风险。

Gartner 建议：在三大支柱上实施 SSCS

软件开发中的监管包括根据已知的安全和运营风险主动评估和批准依赖项和软件工件，以尽量减少未来出现的问题。通过采取主动的方法并考虑各种风险，开发团队可以消除维护不善或不安全的工件，减少仅在项目中包含问题后才发现问题而导致的返工和摩擦。

创造包括两个主要领域：

- 首先，它涉及在开发过程中管理依赖项和工件。虽然尽早评估这些元素（请参阅管理支柱）对于防止引入有问题的依赖项至关重要，但此支柱有助于通过软件物料清单跟踪软件包的使用情况，从而简化未来的漏洞修复。

- 其次，确保开发环境安全至关重要，包括保护开发人员工作站、存储库、版本控制、集成和构建系统以及部署流程。这可以防止对开发基础设施的攻击，近年来，这种攻击已导致严重的供应链漏洞。威胁建模和安全测试等活动（通常与核心应用程序安全活动相关）也通过在代码开发中推广“安全设计”方法增强了供应链安全性。

使用增强了软件采购过程中传统的第三方风险评估。使用由采购和供应商风险管理团队进行，由 SSCS 的技术专家支持，涉及评估供应商软件安全措施并评估交付的代码是否存在策展中讨论的风险类型。此处收集的数据还可以简化对新出现的漏洞的事件响应。安

大事记

齐向东出席 2024 中国互联网大会：以 AI 驱动安全 全力护航互联网新质变

在 7 月 9 日举行的中国互联网协会 2024（第二十三届）中国互联网大会上，全国政协委员、全国工商联副主席、奇安信集团董事长齐向东表示，人工智能新时代，网络安全面临着全方位挑战，要以 AI 驱动安全，应对三大安全威胁、补齐三大薄弱环节，全方位提升新时代安全能力，护航互联网新质变，共同打造安全、稳定、繁荣的网络空间。

齐向东介绍，在“AI 驱动安全”这条路上，奇安信已率先进行了大量探索实践，让 AI 真正赋能安全，实现“质”“效”双提升。一是以 AI 驱动安全，加强分析研判效率；二是以 AI 驱动安全，实现运营能力循环上升；三是以 AI 驱动安全，升级全场景防护能力。



传媒行业再获突破 奇安信中标中新社安全服务项目

近日，奇安信集团中标中国新闻社网络安全及监控服务项目，奇安信将为中国新闻社提供全面的性能检测和网站安全等综合服务。这是奇安信在传媒行业网络安全服务的又一个标杆项目。



齐向东出席 2024 世界人工智能大会：用好“AI+安全”促进数字化转型与可持续发展

7 月 6 日，在 2024 世界人工智能大会（WAIC 2024）AI 重塑通信未来与可持续发展论坛上，全国政协委员、全国工商联副主席、奇安信集团董事长齐向东发表主题演讲。他表示，AI 与安全，是数字化转型与可持续发展的必修课，我国的广大数字经济建设力量，要积极构建高可靠人工智能安全体系，着力为用户提供有效的安全防护，既要用 AI 去驱动安全，又要做好 AI 自身安全，护航数字化转型与可持续发展。

齐向东提出，针对攻击者突破单点设备、突破防护体系、



隐藏踪迹的三个进攻阶段，AI 能给安全防护带来十倍、百倍乃至千倍的效率跃升，将攻击扼杀在事发之前。其一是 AI 提升研判能力，实现安全能力十倍级提升；其二是 AI 提升体系化防御能力，实现安全能力百倍级提升；其三是 AI 提升溯源反制能力，实现安全能力千倍级提升。

陕西洞鉴云侦司法鉴定资质获批准

近日，陕西西安洞鉴云侦声像资料司法鉴定所正式获得陕西省司法厅授予的司法鉴定许可证，成为陕西省近年来唯一获批成立的电子数据司法鉴定机构。这一成就不仅为西北地区的司法鉴定事业注入了新的活力，更标志着奇安信洞鉴的连锁经营模式初见雏形。

奇安信集团和爱数集团达成战略合作 共享数据要素发展机遇

近日，网络安全龙头企业奇安信集团与数据服务领军企业爱数集团签署战略合作协议，双方将在防勒索、安全运营、企业数字化转型战略规划、信创融合、数据安全、数据智能，以及数据交易等方面开展深度产品及解决方案融合，达成全面战略合作伙伴关系，共同引领企业数字化转型发展新方向。



齐向东出席 2024 全球数字经济大会：AI 驱动网络安全“质”“效”双提升

7月2日，2024全球数字经济大会在国家会议中心正式开幕。大会以“开启数智新时代，共享数字新未来”为主题，邀请业内专家、企业领袖和知名学者，就数字经济领域的最新技术趋势、研究成果和应用案例进行分享和探讨。全国政协委员、全国工商联副主席、奇安信集团董事长齐向东受邀出席主论坛，并发表“AI 驱动安全，护航全球数智未来”主题演讲。

他表示，AI 驱动网络安全实现“质”“效”双提升，离不开高质量的数据、体系化的网络安全建设、统一的标准。同时，AI 驱动安全，目标不仅是防范风险，更是要安全释放数字技术的全部潜力，护航全球数智新未来。



响应效率千倍级提升 奇安信发布 AI+SOC 智能安全运营方案

7月3日，2024全球数字经济大会（GDEC）在京举行，众多前沿技术和产品在大会现场首发，展示数字技术创新发展最新成果。奇安信对外发布了 AI+SOC 智能安全运营方案和 AI 代码助手两款新产品。

AI+SOC 智能安全运营方案通过 QAX-GPT 机器人和

NGSOC 产品的深度融合，结合“人工智能模型”和“安全专家经验模型”，可以实现智能分诊、智能研判、智能调查、智能响应等四大功能，旨在解决安全运营工作中海量告警处理难、设备数据联动难、事件响应闭环难等三类难题，驱动安全运营从“密集型”向“智慧型”升级，并实现十倍、百倍、千倍级的效率跃升。

奇安信 AI 代码助手（QAX CodeGen）可提供代码生成、代码补全、代码解释/注释、代码对话等基础功能外，还提供了代码漏洞检查、代码安全修复等聚焦安全场景的专属功能，可以实现开发效率、代码质量和安全性的三重提升，让编程从此更加高效、安全和便捷，助力科技企业释放更大的生产力。

北京市人大专题调研组到奇安信安全中心调研

7月2日，北京市人大专题调研组一行到访奇安信安全中心，调研网络和数据安全保障情况。此行是调研医疗服务信息化便民惠民工作情况中的第三站。与会领导以实例与奇安信安全专家共同探讨了医院的网络安全、数据安全建设。大家一致认为，加强医院的安全生产管理，提高各方面的安全意识，落实安全措施，确保医院的安全和稳定运行势在必行。



奇安信集团与京能集团达成战略合作

6月27日，奇安信集团与京能集团举行战略合作框架

协议签约仪式，双方将在数字安全顶层设计与治理、信创架构安全技术、工控系统网络安全等多领域开展深入合作。京能集团党委书记、董事长姜帆，奇安信集团党委书记、董事长齐向东，副总裁黄翔出席签约仪式。奇安信集团副总裁张龙与京能集团副总经理王永亮代表双方签署协议。



奇安信集团与交通运输部公路院签订战略合作框架协议

6月20日，奇安信集团与交通运输部公路科学研究院（以下简称“公路院”）举行了战略合作框架协议签约仪式。双



方秉承“互惠互利、优势互补、长期合作、共同发展”的原则，将充分发挥技术和资源优势，围绕网络安全能力建设、数据安全能力建设、密码应用安全研究、行业网络安全技术、交通强国试点攻关、行业安全产品孵化等方面，以行业网络安全技术自主可控为目标，在公路交通运输数、信、智、网安全技术研发和创新应用领域展开全面深入合作。

齐向东受聘北京交通大学兼职教授仪式暨名师大讲坛活动圆满举办

6月20日下午，齐向东兼职教授聘任仪式暨名师大讲坛活动圆满举办。全国政协委员、全国工商联副主席、奇安信集团董事长齐向东受聘为北京交通大学兼职教授。聘任仪式上，关忠良副校长为齐向东颁发聘书及校徽。

聘任仪式后，齐向东作题为“网络安全新形势下的应对举措和发展策略”的学术报告。报告结合奇安信集团在网络安全方面的调查研究，从“网络安全问题源于漏洞不可避免”“常态化网络攻击成为新形势”“人工智能带来的四大安全挑战”“用内生安全筑牢安全‘底板’”“加速网安产业与人才双向赋能”五个方面进行分享，通过数据的分析和案例的讲解，帮助同学们深刻体会到网络安全的重要性及人工智能对网络安全的深刻影响，“人工智能+”为我国迈向数字经济强国提供宝贵机遇的同时，也带来难以回避的安全风险。



奇安信集团与中国电信陕西公司达成战略合作

近日，奇安信集团与中国电信陕西公司在西安签署战略合作协议。双方将在内生安全、政企 ICT 领域及联合创新和行业发声等方向进行全面的战略合作，共同打造战略协同、优势互补、资源共享、共赢发展的安全服务业务生态链，促进双方在各自领域创造更大的经济效益和社会效益，开创产业发展新格局。



奇安信集团与中保车服签署战略合作协议

6月18日，奇安信集团与中保车服科技服务股份有限公司（以下简称“中保车服”）举行了战略合作协议签约仪式。



式。双方将充分发挥各自优势，围绕“一室、两标准、三中心”开展合作，共同探索金融科技创新的新模式。中保车服总裁王小明、总裁助理兼 CTO 王晓斌，奇安信集团高级副总裁吴登峰、副总裁唐皓出席签约仪式。

奇安信集团 77 个项目获批教育部第三期供需对接就业育人项目

近日，教育部高校学生司发布《关于公布教育部第三期供需对接就业育人项目立项名单的通知》（教就业司函〔2024〕23 号）。按照工作流程，经高校与用人单位联合申报、专家审核，奇安信科技集团股份有限公司联合高校申报的 77 个项目成功立项，其中就业实习基地项目 24 个，定向人才培养培训项目 42 个，人力资源提升项目 11 个。

奇安信与金山办公达成战略合作 携手共筑数字化办公安全防线

6 月 13 日，奇安信集团与金山办公在北京签署战略合作。双方将在网络安全技术创新、产品开发及市场拓展等多个维度展开深度合作，共同推动我国办公软件行业的安全升级与智能化转型。

此外，双方还计划共同打造联合创新实验室，开发一系



列覆盖终端安全、边界安全、漏洞与威胁情报、AI 安全、代码安全、数据安全、云安全、信创安全等多个层面的安全解决方案。这些方案将为各种数字化办公场景提供强大的威胁防护能力，确保企业数据和运营的安全性。

齐向东出席市政协议政会 建言助力北京新型工业化发展

6 月 3 日，市政协和市委统战部联合召开议政会，围绕“加快先进制造业与数字经济有机融合，推进新型工业化北京实践”协商议政。市委书记尹力出席会议并讲话，市政协主席魏小东主持。17 位民主党派市委、市工商联和无党派人士代表及政协委员、专家学者发言，围绕中试数字化智能化建设、发展高级别自动驾驶、提高生物医药大数据共享能力、加大数字关键核心技术攻关力度、筑牢数字基础设施底座等建言资政。

全国政协委员、全国工商联副主席、政协委员、市商会副会长、奇安信集团董事长齐向东作为市工商联代表，做“筑牢数字基础设施底座 助力北京新型工业化发展”主题发言。



新质生产力专题讲座成功举办 齐向东授课

为进一步推动国企改革深化提升，提高国有企业核心竞

争力，促进扬州市国资国企新质生产力培育发展，5月28日下午，由扬州国资委主办、科教集团（大数据集团）承办的第35期国资大讲堂在扬州花园国际大酒店成功举办。

此次讲座邀请了全国工商联副主席、全国政协委员、奇安信集团董事长齐向东，以《统筹好发展和安全两件大事 实现新质生产力跃升》为题进行专题授课。齐向东深刻阐述了习近平总书记有关“新质生产力”的重要论述，强调了网络安全护航新质生产力的重要性，并就统筹好发展和安全两件大事提出了意见建议。



齐向东：发展为先、安全护航 推动数据要素发挥乘数效应

5月24日，以“释放数据要素价值 发展新质生产力”为主题的第七届数字中国建设峰会在福州开幕。在“数据资



源与数字安全”分论坛上，齐向东表示，高质量的数据是驱动经济社会高质量发展的核心要素，要坚持“发展为先、安全护航”，以促发展为终极目标，推动数据要素高效流通。

他强调，要以促发展为最终目标，整体性构建数据安全合规保障体系，让数据要素更好地应用在工业、农业、金融、科技创新等重点领域发挥出更大的乘数效应。一是要确保数据基础设施安全，构建纵深防御的内生安全体系；二是确保数据流转全过程安全，从业务视角出发做好全局防护、重点防护；三是确保数据要素化全过程合规，构建数据流通安全合规体系；四是确保数据安全防护有实效，通过持续安全运营提升安全能力。

齐向东：“人工智能+”赋能“数据要素x”安全释放数据要素价值

5月23日，由中国人工智能学会主办、福建省大数据集团有限公司承办的“人工智能与数据要素产业生态大会”在福州举行。奇安信集团董事长齐向东受邀出席并做主题演讲。

他表示，“人工智能”与“数据要素”就像是发动机与燃料，是引领数字经济发展的黄金组合，但也不可避免带来新的挑战，亟需打造多层次的合规保障体系，促进数据合规、高效的流通使用，安全地释放数据要素的乘数效应。一是以纵深防御的内生安全体系，从底层筑牢数据要素流通基础设施；二是以AI安全整体解决方案，从中层加强数据要素流通



平台安全保障；三是以“可用不可见”的安全技术，从顶层构建可信的数据要素流通环境。

豫信电子科技集团董事长李亚东一行到访奇安信

5月20日，豫信电子科技集团党委书记、董事长李亚东一行到访奇安信安全中心，与奇安信集团党委书记、董事长齐向东进行会谈，双方就推动河南产业数字化转型、加强数据要素产业生态体系构建、深化数字政务数据安全保障等内容进行深入交流。



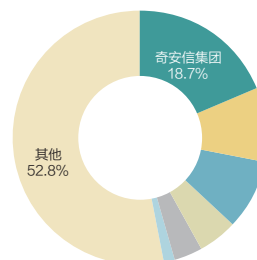
荣誉墙

赛迪报告：奇安信终端、安管、安服再度蝉联市场第一

赛迪顾问近日发布的《2023—2024 年中国网络信息安全市场研究年度报告》指出，2023 年，中国网络信息安全整体市场在政企用户普遍降低或者延后相关预算后，依然保持较缓增长，规模达到 998.3 亿元，同比增长 8.5%。到 2026 年，网络安全市场规模将为 1358.6 亿元，三年复合增长率达到 10.8%。

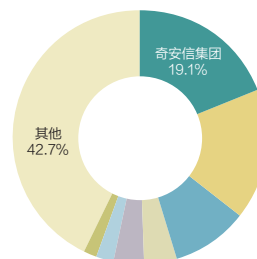
报告显示，奇安信集团在 2023 年以 64.4 亿元的营业收入、6.5% 的市场销售总收入占比位居行业第一。其中，在安全软件的终端安全、安全管理平台领域，以及安全服务中，奇安信分别以 18.7%、19.1%、8.4% 市场占有率排名第一，连续多年处于行业领导者；传统硬件市场中，UTM（统一威胁管理）、Web 安全产品销售份额排名前二，拥有稳固的市场地位。

2023 年中国终端安全市场厂商结构



数据来源：赛迪顾问 2024，02

2023 年中国安全管理平台市场厂商结构



数据来源：赛迪顾问 2024，02

2023 年中国安全服务市场厂商结构



数据来源：赛迪顾问 2024，02

首批、首家！奇安信椒图通过可信安全“云上勒索攻击防护”检验

近日，奇安信网神云锁服务器安全管理系统（椒图）在“云上勒索攻击防护准备、事件预防、事件检测、事件缓解、事件恢复、根源分析”六类指标的能力方面，率先通过了 Q/KXY CS118——2024《云上勒索攻击防护第 2 部分：系统安全》标准的检验，在国内首批、首家获得中国信通院、泰尔实验室联合颁发的《云上勒索攻击防护系统安全能力检验证书》。



全领域覆盖 奇安信入选《数字安全护航技术能力全景图》

日前，中国信息通信研究院“数字安全护航计划”正式发布首期《数字安全护航技术能力全景图》。奇安信凭借全面、卓越的技术和能力，成功入选全景图全部 14 大安全领域、103 个细分领域。彰显了其在网络安全领域的全面覆盖能力和行业领先地位。

在激烈的竞争中，奇安信从 391 家申请者中脱颖而出，经过评估，最终有 147 家厂被录入首期全景图，这不仅是对奇安信技术实力的认可，也是对其在数字安全领域贡献的肯定。奇安信将继续致力于技术创新，为数字中国的建设贡献力量，守护国家和企业的重要信息资产，助力数字经济的稳健发展。



入选 Gartner®《2024 中国基础设施战略成熟度曲线》报告 被列为“代表性供应商”

近日，Gartner 发布了《Hype Cycle™ for Infrastructure Strategies in China, 2024》报告（《2024 中国基础设施战略成熟度曲线》，以下简称《报告》）。分析了当前中国市场企业组织的基础设施战略面临的挑战，并给出了针对性的投资建议。其中，奇安信被列为零信任网络访问（ZTNA）领域代表性供应商（Sample Vendors）之一。

作为国内零信任安全的首倡者，奇安信多年来持续专注

“零信任安全架构”研究与发展，始终积极推进零信任落地与企业 IT 架构改革。目前，奇安信在零信任体系的建设与落地方面积累了丰富的经验。近日，由奇安信牵头的我国首个零信任国家标准《网络安全技术 零信任参考体系架构》（GB/T 43696-2024）也正式发布。

奇安信获评“2023 年度漏洞报送最具贡献单位”

近日，工业和信息化部网络安全威胁和漏洞信息共享平台（NVDB）对 2023 年在网络产品安全漏洞管理方面做出杰出贡献的单位进行表彰。凭借在漏洞报送数量和质量方面的突出表现，奇安信荣获“2023 年度漏洞报送最具贡献单位”。

自 2019 年 NVDB 上线以来，奇安信一直积极参与并全力支持平台建设，特别是在通用网络产品安全漏洞专业库方面，贡献显著。此前，凭借对 NVDB 五大漏洞专业库的全方位支持，奇安信获评“2022 年度漏洞治理合作最具贡献单位”。



奇安信企业级防火墙再获国际权威机构推荐

日前，奇安信企业级防火墙凭借卓越的产品技术和安

全解决方案实力，强势入围 Forrester《The Enterprise Firewall Landscape, Q2 2024》报告，获得全球代表性供应商推荐，这也是奇安信智慧防火墙连续三次获得国际权威机构 Forrester 的认可和推荐。

报告对全球 17 家主流企业级防火墙供应商的市场规模、产品构成、业务焦点进行展现与分析，以期为客户提供更符合自身需求的参考维度。奇安信企业级防火墙继入选 Forrester《Now Tech: Enterprise Firewalls, Q1 2020》、《Now Tech: Enterprise Firewalls, Q2 2022》之后，再次获得该项目的推荐。

奇安信再次荣登“科创中国”先导技术榜

7 月 2 日，中国科学技术协会（以下简称“中国科协”）发布 2023 年“科创中国”系列榜单遴选结果，奇安信集团“基于应用程序接口精准检测和高效动态防护的云安全关键技术与应用”项目在众多项目中脱颖而出，作为网络安全领域唯一代表性技术成果入选“科创中国”先导技术榜。奇安信集团已经连续两年入围“科创中国”先导技术榜。

A blue poster for the 2023 "Science and Innovation China" series ranking. It features the title "2023年‘科创中国’系列榜单" and "先导技术榜——电子信息领域". Below the title is a table listing four technologies and their recommending organizations.

第二十六届中国科协年会			
2023年“科创中国”系列榜单			
先导技术榜——电子信息领域			
(按推荐单位排序)			
序号	技术名称	技术团队	推荐单位
1	使用大型柔性太阳翼的平板堆叠式卫星——银河航天天翼星 03 星	银河航天（北京）网络技术有限公司	中国电子学会
2	超长距离超低时延开放解耦的高速率传输技术	上海欣诺通信技术股份有限公司	中国通信学会
3	基于亿阻器新型器件的存算一体关键技术	中国移动通信有限公司研究院	中国通信学会
4	基于应用程序接口精准检测和高效动态防护的云安全关键技术与应用	奇安信科技集团股份有限公司	中国通信学会

连续三年排名第一 奇安信领跑 CWPP 云工作负载保护平台市场

近日，全球领先的 IT 市场研究和咨询公司 IDC 发布了《中国私有云云工作负载安全市场份额，2023：CNAPP 助力企业实现全方位云原生安全防护》，报告指出，中国私有云云工作负载安全市场在 2023 年实现了 3.5% 的同比增长，市场规模达到 15.3 亿元人民币。

其中，奇安信 2023 年在政府、金融、运营商、卫生、能源、制造等多个重要行业取得连续突破，帮助客户建成了多个重点行业的云工作负载安全示范级标杆，以 25.8% 的占有率连续三年排名第一，领先优势进一步扩大。

2023年中国私有云云工作负载安全（CWPP）市场份额



数据来源：IDC《中国私有云云工作负载安全市场份额，2023：CNAPP助力企业实现全方位云原生安全防护》报告

奇安信上榜 Gartner® 安全信息与事件管理魔力象限

近日，全球 IT 咨询机构 Gartner 发布了 2024 年 SIEM 魔力象限（Magic Quadrant™ for Security Information and Event Management）报告，即业界所熟知的安全信息和事件管理魔力象限报告，俗称 SIEM 产品的“试金石”。奇安信本次上榜 Gartner SIEM 魔力象限，意味着 NGSOC 进入国际领先产品行列。

据悉，本次报告直到今年 5 月份才正式发布，这距离上一本 SIEM 魔力象限报告已相距 1 年零 7 个月之久。在激烈的竞争中，奇安信赫然上榜，这不仅是厚积薄发的体现，更是奇安信 NGSOC 迈向全球领先战略旗开得胜的标识。

奇安信荣获“长城杯”优秀技术支撑单位 助推高校实战型人才培养

第一届“长城杯”信息安全铁人三项赛决赛于上月圆满落幕，奇安信集团作为赛事重要承办单位，荣获由中国信息安全测评中心颁发的“优秀技术支撑单位”称号，充分展现了奇安信全面卓越的竞赛支撑保障实力，为我国打造实战型网络安全专业人才贡献力量。

作为大赛重要的承办单位，奇安信深度参与了题目设置、裁决判定、技术支撑、运营保障等各个环节：题库方面，奇安信凭借深厚的行业积累和技术洞察，设计了一系列既考验参赛者理论知识又注重实战技能的题目，并对题库难度和知识进行审核；裁判环节，奇安信资深安全专家以公正、严谨的态度进行评判，保证了比赛结果的公平性和权威性。



奇安信获得国家科学技术进步奖

6 月 24 日，2023 年度国家科学技术奖在京揭晓，奇安信科技集团股份有限公司参与申报的“超大规模多领域融合联邦靶场（鹏城网络靶场）关键技术及系统”项目获得国家科学技术进步二等奖。这也是本年度网络攻防领域唯一的国家级科学技术进步奖。

此次获奖是奇安信强研发、重创新战略所取得成绩的重要体现。截至目前，奇安信承担国家重大专项、示范工程共计 94

项，攻克一批网络安全核心技术，荣获省部级科技进步相关奖项 26 项。参与国家标准制定 75 项，其中牵头 2 项，已完成 49 项，是国内参加网络安全领域国家标准最多的安全企业。



奇安信四项信创解决方案全部上榜“2024 广东软件风云榜”

日前，由羊城晚报报业集团联合广东软件行业协会、广东省大数据协会征集评选的《2024 年广东软件风云榜》，在 2024 年粤港澳软件产业高质量发展大会、第十二届粤港澳云计算大会暨第七届粤港澳 ICT 大会上正式发布。榜单表彰为广东软件产业和数字经济、新型工业化发展作出突出贡献的企业、企业家、技术骨干、优秀产品等。

奇安信终端安全一体化解决方案、奇安信云安全管理平台解决方案、奇安信网神跨网文件安全交换管理解决方案，获评“优秀信息技术应用创新行业应用解决方案”；奇安信态势感知与安全运营解决方案，摘得“信息技术应用创新行业解决方案 TOP 10”。

再获第一！奇安信天眼连续三年领跑国内 NDR 市场

IDC 最新发布的《中国网络威胁检测与响应市场份额，2023：NDR 正在向更多技术栈延伸》（Doc # CHC50964624，2024 年 6 月）显示，2023 年中国网络威胁检测与响应市场较 2022 年同比增长 1.6%，规模达 23.8 亿元人民币，其中奇安信市场份额遥遥领先，占比达到 22.5%，连续三年排名第一。

目前，奇安信天眼（威胁监测与分析系统）已经为国内不同行业的 6000 家政企客户提供了集监测预警、威胁检测、溯源分析、响应处置为一体的网络威胁监测与分析解决方案，发现和处置超过百起 APT 攻击事件，赢得了客户的广泛信赖，树立了国内 NDR 产品领域的标杆形象。

2023年中国网络威胁检测与响应（NDR）市场份额



数据来源：IDC《中国网络威胁检测与响应市场份额，2023：NDR 正在向更多技术栈延伸》报告

综合排名第一！奇安信揽获 8 项 CNNVD 重磅大奖

6 月 18 日，国家信息安全漏洞库（CNNVD）2023 年

度工作总结暨优秀支撑单位表彰大会在中国信息安全测评中心隆重举行。

奇安信作为 CNNVD 一级技术支撑单位，凭借在漏洞挖掘、漏洞消控、漏洞通报等方面的突出贡献，以综合排名第一名的成绩在 262 家技术支撑单位中独占鳌头，荣获 8 项重磅大奖，分别为“优秀技术支撑单位”“高质量漏洞优秀贡献单位”“漏洞消控优秀贡献单位”“高质量通报优秀贡献单位”“优秀特聘技术专家”，以及 3 项“漏洞奖励一级贡献奖”，其中“漏洞奖励一级贡献奖”获奖数量位居第一。



《云安全资源池能力指南》发布 奇安信荣获创新力及市场执行力双第一

近日，国内某知名研究机构发布《云安全资源池能力指南》，报告指出随着云计算的迅速发展，越来越多的政企用

户选择业务上云，与此同时，云上数据和应用将面临着各种安全威胁，如数据泄露、病毒感染、网络攻击等，能够集中统一管理各类云安全资源的云安全资源池应运而生，目前国内云安全资源池产品快速发展，市场规模已经达到 12.8 亿元人民币。奇安信云安全资源池凭借出色的产品性能和市场占有率，获得创新力及市场执行力双第一。



奇安信荣获网络安全国家标准优秀实践案例一等奖

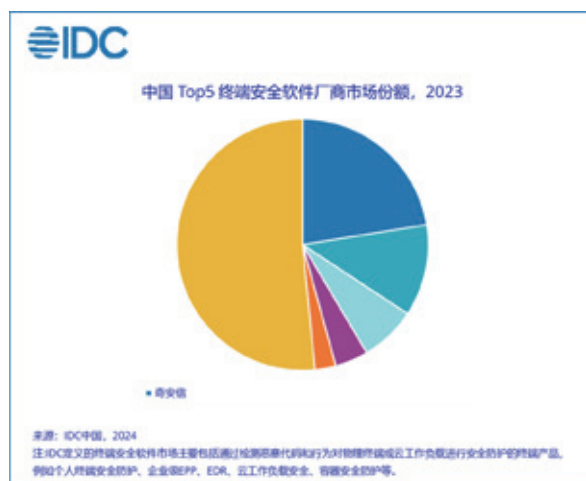
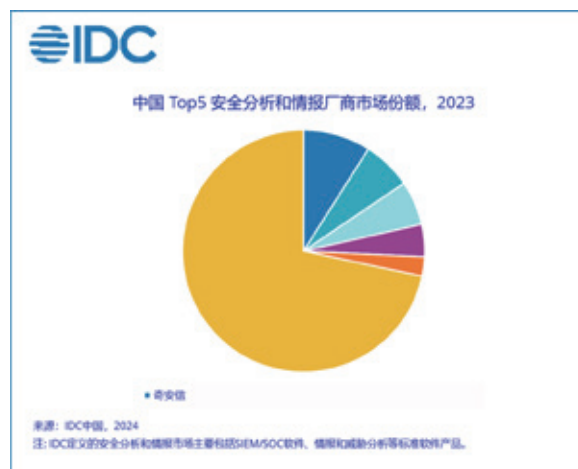
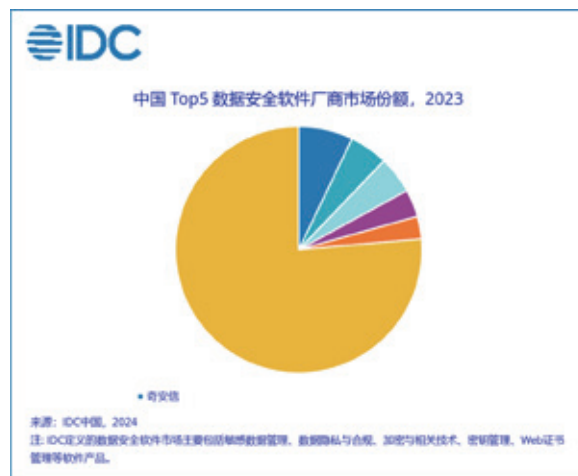
近日，全国网络安全标准化技术委员会陆续向 2023 年网络安全国家标准优秀实践案例获奖单位发放荣誉证书。由奇安信参与和牵头两大案例分别荣获一等奖和二等奖。





三大市场持续领先！奇安信终端安全、数据安全、分析和情报市场再获第一

5月13日，全球领先的IT市场研究和咨询公司IDC发布了《中国IT安全软件市场跟踪报告，2023H2》（以下简称《报告》）。《报告》统计显示，奇安信凭借卓越的技术实力和优秀的市场表现，再次领跑终端安全、数据安全、安全分析与情报市场。特别是终端安全领域，在奇安信连续六年稳居市场第一的基础上，再次扩大市场占有率，彰显了强劲的产品竞争力和市场地位。



社会责任

“眼明心安”项目走进西藏亚东 为边境县带起一支眼病筛查队伍

6月30日—7月5日期间，北京白求恩公益基金会、北京奇安信公益基金会、北大医学眼科专家共同发起的“眼明心安—西藏儿童盲及低视力诊疗能力提升项目”（以下简称“眼明心安”项目）2024公益筛查带教活动来到了祖国西南边陲——日喀则市亚东县，开展儿童眼病筛查带教工作。

筛查团队用5天时间，对亚东县城小学和幼儿园的453名学生、吉汝乡中心小学的114名学生、堆纳乡中心小学的308名学生，以及帕里镇小学的407名学生，共1282名学生进行了眼病筛查，并为122名亚东群众进行了义诊。同时，全县七个乡镇的五官科、全科、藏医科等科室精心挑选的10名骨干医生跟随筛查团队进行了全程学习。经过实践与考核，10名当地医生都可以独立胜任儿童眼健康及眼病筛查所涉及的各项检查检测，具备了基本的常见眼病的诊疗能力。



“眼明心安”项目荣获“520 社会责任日”关爱儿童议题优秀案例

6月28日，2024年“520 社会责任日仲夏夜”在北京成功举办，上百位来自企业、基金会、行业协会、学术机构和媒体机构的代表出席，共同见证了“520 社会责任日案例”评选结果的最终揭晓。北京奇安信公益基金会“眼明心安·西藏儿童盲及低视力诊疗能力提升项目”（简称“眼明心安”项目）荣获关爱儿童议题“优秀案例”。

奇安信基金会相关负责人表示，此次获奖是对“眼明心安”项目的充分肯定。项目组将继续通过多方协作和有效整合资源，将优质医疗资源引入西藏，并利用好医疗人才“组团式”援藏平台，结合实际情况，探索出切实有效的儿童眼健康筛查、防治和诊疗的公益模式，使得西藏自治区的孩子

们都享有看得见、看得清的权利，拥有更加“光明”的未来。



“眼明心安”项目走进西藏江达

6月15日，“眼明心安”项目2024公益筛查带教活动正式拉开帷幕。第一批儿童眼病筛查团队，携带2台视力筛查仪和4台眼底照相机及协调的2台蔡司生物测量仪器，从北京出发抵达西藏自治区昌都市江达县。5天时间里，筛查团队为来自江达县三所幼儿园、两所小学，以及两所初中的约1400名学生进行了眼病筛查，每天平均筛查280余人。

此外，本次筛查活动也强化了当地医护人员对眼健康重要性的认识，使他们更加明白对儿童眼睛进行早筛的必要性。这5天的筛查不仅是一次技术学习的机会，更是一次全面提升眼科诊疗服务质量的实践。江达县人民医院目前尚未设置



眼科，参与本次筛查活动的本地医护们，有望成为未来眼科建设的储备人才。

“眼明心安” 2024 公益筛查带教活动出征仪式在京举行

6月12日，由北京白求恩公益基金会、北京奇安信公益基金会、北京大学人民医院眼科专家等共同发起的“眼明心安—西藏儿童盲及低视力诊疗能力提升项目”2024公益筛查带教活动出征仪式在北京奇安信安全中心举行。

据悉，今年“眼明心安”项目将通过捐赠筛查设备、西藏儿童眼病筛查、持续督导基层眼科医生等多种手段，重点聚焦提升西藏眼科医生技术水平与诊疗能力，同时，继续推进理论与实操培训、诊疗与手术带教，为西藏眼科医生培养和儿童眼健康事业做出更大贡献。



“心安助学·高校教育助学”项目举行师资培训班

5月31日，由北京奇安信公益基金会发起的2024年度“心安助学·高校教育助学”项目正式启动，并在湖北大学举行了首期师资培训班，为本年度19所项目院校的42名教师进行了为期3天的主题培训。

本次师资培训班面向2024年度19所项目院校教师展开，由奇安信高校合作中心、补天漏洞响应平台相关专家担任讲师，围绕“网络安全竞赛模式讲解及实操”“网络漏洞挖掘实战能力提升”两大主题进行培训。

奇安信基金会相关负责人表示，接下来受训老师将会回到各自学校，承担起项目负责人角色，负责组建和指导学生

社团，为后续参与“心安助学·高校教育助学”项目实践活动进行准备。



2024 年度心安助学项目启动 19 所资助高校名单发布

5月31日，北京奇安信公益基金会“心安助学·高校教育助学项目”2024年项目启动会在湖北大学会议中心正式举行。会上公布了“心安助学·高校教育助学项目”的全新定位及发展目标，公布了本年度19所资助高校名单。

在此前的学生救助和激励基础上，从今年起，项目将重点围绕教师培养和社会实践支持两个方面开展工作。在教师培养方面，将每年举办网络安全教师公益培训班，为资助院校的教师提供相关专题培训，提升老师在网络安全领域包括漏洞挖掘、渗透测试、安全防护等实战能力和指导能力；在社会实践方面，将支持网络安全相关学生社团的组建，并提供相关赛事、实战保障等机会，为学生提前接触实战化环境提供平台和机会，为接触实际岗位需求搭建实践桥梁。





终端 BG- 国产化事业部 王洪飞摄 地点：金山岭长城



架构演进叠加客户需求，推动边界安全加速融合

作者 吴亚东

IDC 不久前发布数据显示，2023 年，中国网络安全硬件产品市场规模达到 225 亿人民币，到 2027 年，中国网络安全硬件市场规模将达到 364 亿人民币。这其中，边界安全是占比最大的主力产品，也是广大政企客户安全建设的刚需，长期占据着网络安全投入的主要比重。

以防火墙为代表的边界安全，究竟是如何成为行业“常青树”的？其发展 30 多年来经久不衰的密码又是什么？我们的答案是：边界安全的发展史就是能力的融合史。

并非巧合，两大巨头都因“融合”而起飞

2023 年，全球收入排名前二的两家网络安全公司分别是 68.93 亿美元的 Palo Alto Networks（派拓网络），以及 53.05 亿美元的 Fortinet（飞塔），巧合的是，这两家公司从创立到现在，均走的是一条能力融合的路线。

Fortinet 公司从 2000 年创立之初，就打造了全新的概念——UTM（统一威胁管理）。UTM 防火墙是一种集成多种网络安全功能于一身的设备，常见功能包括防病毒、防垃圾邮件、入侵检测与防御、VPN、内容过滤等，其核心是把所有功能融合到一台防火墙上，这是边界安全在融合演进路线上的重要里程碑之一。

然而，UTM 的融合之路并不平坦，由于当时软件技术架构不够先进，使用各种模块的串联检测效率低下，性能消耗大，尤其是当时 CPU 多核技术还没有出现，软件理念的发展和硬件的发展无法匹配，所以 UTM 在发展之初，被诟病的主要是功能融合后带来的性能下降问题。

2005 年，一家名为 Palo Alto 的公司诞生，该公司在融合的路线上继续迈进一大步，并和 Gartner，联合提出下一代防火墙（NGFW）的概念，它基于一体化引擎和多核硬件并行处理，实现多种安全功能。由于在



图：奇安信集团副总裁、首席架构师吴亚东

这个时期，CPU 的多核也发生巨大的变化，性能得以快速提升，下一代防火墙可谓生逢其时，很快风生水起。

NGFW 不仅具有高可用性、NAT、ACL、VPN 等基础安全接入能力，还融合了应用识别、URL、IPS、AV、WAF、SDWAN、ZTNA、物联网、协同联动等高级安全能力，很快成为市场主流。

防火墙的演进是边界安全发展的缩影。事实上，在 UTM 之前，20 世纪 90 年代的传统防火墙本身也是融合的理念，因为其前身是具备简单 ACL 功能（访问控制列表）的路由器，通过融入基础安全接入、高可用性、NAT、ACL、VPN 等功能，加上硬件架构的变革，形成专有的防火墙产品。从 90 年代至今的 30 多年中，以防火墙为代表的边界安全，始终是规模第一的网络安全细分赛道。而沿着能力融合路线的 Fortinet 和 Palo Alto Networks，其市值也分别达到了 400 多亿美元和 1000 亿美元左右。

架构演进叠加客户需求，推动边界安全加速融合

从 UTM 到 NGFW 的演进，不难发现，能力融合的关键是需要硬件和软件架构的更新，两者缺一不可。因为在单核时代，设备在处理多个任务时性能就会急剧下降，无法融合更多的功能。而多核甚至众核时代的到来，边界安全的融合演进，才得以大大提速，推动了防火墙第四代软件架构——融合网关的出现。

在硬件层面，以飞腾为代表的国产多核高性能 CPU 正在发展，其中服务器芯片腾云 S5000C 最高支持 64 核，是国内首个采用 DDR5 和

以防火墙为代表的边界安全，
究竟是如何成为行业“常青树”的？
其发展 30 多年来经久不衰的密码又是什么？
我们的答案是：
边界安全的发展史就是能力的融合史。

PCIe5.0 的芯片，支持两路直连，32 核支持 4 路互联，这为融合网关支持更多的功能集成，提供了底层算力基础。

从技术演进的角度来看，一个产品要融合各种能力，其实有两种做法。一种是功能组件级的融合，即将各种模块作为内部的一个功能组件融合到这个网关里面，代表产品为融合网关。另一种是基于虚拟化的融合，其代表产品是 uCPE，目前在国外比较普遍，其缺点是无法支持硬件交付。

不仅是安全能力的融合，简化边界更是融合的长期目标。因为传统的网络边界防御场景通常是以“串糖葫芦”的方式，依次将各种网络安全设备进行串联，为了增加可靠性，一般还会进行冗余的方式工作，这种部署方式存在以下问题：

首先是故障率高，且缺乏弹性，维护困难。比如，四个安全设备串在一起，一个设备的可靠性是 99%，那 4 个设备的安全性就是 99% 的 4 次方，急剧下降至 96%。同时由于菊花

链式的多级串接，故障点多，排障复杂，同时设备软 / 硬件变更会造成较长时间业务中断。当 PoC 测试时只能全流量接入，业务影响面大，且多厂商无法并行。

其次是性能浪费，工作效能低。串联中的设备不得不处理很多无关数据，无法仅接收擅长处理的流量，造成大量资源浪费，工作效能低，客户 TCO 居高不下。同时，任何安全设备性能瓶颈将影响整个业务路径，很容易导致业务受损。

再次是扩展性差，容易影响业务。在业务扩容需要更换更高性能的设备时，替换割接需要断网；在安全能力需要扩容时，新增安全设备上线也需要断网割接，会影响客户业务正常运作。

最后是解密性能不足，造成安全隐患。当所有设备开启 SSL 解密时，性能衰减严重。多设备同时解密，延迟倍增，同时证书管理困难。甚至大部分解密流量没有被解密，安全检测效果受到极大影响，给隐藏的加密攻

击留下了可乘之机。

为了解决这些难题，奇安信推出了基于融合理念的流量解密编排器（SSLO），它通过重构安全设备部署架构、安全设备资源池化、业务按需进行编排引流、SSL 解密流量编排等技术手段，实现了一种弹性、资源池化、大加密流量场景可检测、管理更及时及响应更敏捷的集成方案，让网络维护化繁为简，实现了安全、高性能、易管理的完美平衡。

端网融合、数通融合，跨领域融合将成为未来趋势

展望未来，融合网关将不仅仅是边界安全能力融合，还会将更多安全之外的跨领域能力融合进来。其中包括端网融合，即终端安全和边界安全的融合及数通融合，将数据通信和边界安全功能融合进来。

奇安天合安全融合机是终端和边界安全能力融合的代表产品。其核心是将终端的控制台直接融合到出口网

关之中，成为一个极简方案。一个出口的网关既能做边界安全使用，同时又能够作为终端的控制台使用，这样合体之后，实现了三种能力的融合：管理融合、防御融合、感知融合。

举个例子，过去在出口网关上通过一个威胁情报发现某个事件，安全人员会知道包括 IP、威胁情报 ID 等信息。但用户会进一步发出疑问，比如，这个到底有没有出问题？如何研判？如何处置？通过将终端和边界做了融合以后，除了 IP 相关信息，还能够知道这个事件的所有信息，包括操作系统的详细版本信息，是否有相关漏洞，是否需要打补丁，包括对这个终端上进行全盘查杀，查找相应的威胁情报对应信息等。这样通过端和网的打通，就能进行更全面的数据挖掘和发现，避免单个情报导致的误报等。可以说，这是一个完全实现端网融合的产品，而不是简单联动的方案。

除了端网融合，数通和安全也可以实现融合。通过部署“安全和网络融合”的安全网关，可实现有线和无

线网络及设备的统一的安全性、连接性和访问控制管理，从而企业可轻松地跨有线和无线集中管理一致的用户、设备和应用程序策略。

在实现原理上，融合网关和接入设备组成一台大的逻辑网关，交换机端口和 AP 都是安全网关的延伸端口，构建一个逻辑的纵向堆叠安全设备。所有流量经过安全网关绕行以进行安全流量管控，其中东西向流量可以通过二层 GRE 或 VxLAN 隧道方式打到网关进行安全处理。同时，融合网关可收集整网的时候拓扑，通过轻量级管控界面实现网络和设备的统一管理、安全分析和联动处置，这无疑大大简化了 IT 人员的运维操作。

不止边界，融合代表着整个行业创新方向

当前，融合、平台化集成逐步成为全球网络安全行业的共识。早些年，Gartner 曾提出“网络安全网格架构（CSMA）”的理念，其核心是为客户提供一体化的安全结构和态势。今年 RSAC 大会，越来越多的厂商推出了集成解决方案，将多个产品和服务结合到同一安全架构中，为客户提供更好的协同性和融合性。可见，融合不仅是边界安全的发展趋势，更是整个网络安全行业的创新方向。正如国家信息中心办公室副主任吕欣在 2024 北京网络安全大会（BCS 2024）“融合安全论坛”上所说的，“在网络系统的跨层级、跨地域、跨系统、跨部门、跨业务的协同管理和服务越来越普遍的情况下，技术融合、数据融合、业务融合成为趋势，推动网安行业尤其是融合安全技术创新和升级迭代变得更为迫切。”

关于作者

吴亚东

奇安信集团副总裁、首席架构师



人工智能对密码学的威胁探讨

作者 洪璐 方婷

随着人工智能技术的不断发展，其与密码学之间的联系愈加紧密，彼此之间的相互作用也愈发深入。近年来，有关人工智能可能对密码学带来的影响引起了越来越多的讨论。一方面，人工智能在帮助密码学提高安全性和效率，同时又能利用密码学保护自身的数据，与密码学形成相辅相成的发展态势；另一方面，人工智能的发展（尤其是在机器学习和深度学习方面）也增强了密码分析的能力和效率，并因此危及到密码系统的安全性，为密码学带来了一系列潜在的威胁。

本期简报对 ResearchGate 上的一篇系统性文献综述的关键内容进行了编译和总结，该篇综述由来自斯里兰卡萨伯勒格穆沃大学的 Pinindu Chethiya 发表于 2023 年 12 月。他搜集了大量探讨相关议题的文献，总结并批判性地评估了有关人工智能对传统加密方式影响的研究现状，讨论了该领域当前研究存在的局限性和面临的挑战，并对未来的研究提出了建议。

一、介绍

随着人工智能及相关技术的高速发展，人工智能算法已经具备了发现加密系统中脆弱性与漏洞的能力，这使得加密数据的安全性受到了威胁。近年来，人工智能技术已被应用于各种密码攻击活动中，如破解密码、密

钥预测、绕过安全协议等。与此同时，有关人工智能对密码学威胁的研究和观点也层出不穷。因此，了解与加密系统相关的潜在风险和漏洞，并开发新的方法来防止基于人工智能的攻击已经势在必行。

本文综述旨在总结人工智能对密码学影响的研究现状，概述人工智能对密码学构成的潜在风险和威胁，并确定该领域的主要趋势和未来研究方向。综述内容涉及不同类型的基于人工智能的密码攻击、已经提出的防御对策，以及这些对策存在的挑战和局限性。本综述还分析了人工智能和密码学的最新进展及对密码学未来的潜在影响，包括其对隐私、安全和机密性产生影响的潜在后果。

二、人工神经网络 (ANNs)、机器学习和深度学习

目前的许多观点和研究中有相当一部分聚焦于人工智能在国防领域的运用，然而，这其中有不少观点实际上已经不再适用于高速发展中的人工智能领域，原因就在于这些观点并没有敏锐和密切地关注到人工智能在密码学中的应用。

与此同时，另有一些研究者已经关注到了人工智能在密码学中的应用，他们的研究主要聚焦在使用人工智能辅助密码学的几个特定环境，其中就

包括对人工神经网络、机器学习和深度学习等人工智能技术的使用。近来的许多相关研究都发现了这几项人工智能技术对入侵检测和身份验证的助益，并已将其与其他技术相结合后实际应用于多个领域，例如，区块链技术和人工智能结合后在智能汽车的身份验证及加密方面的应用、深度学习在车与车（V2V）交互中的应用，以及机器学习在智能电网中的应用等。也有一些研究提到了人工智能技术在应对网络攻击中的应用，例如，在应对 DDoS、窃听、木马植入等各种边缘网络攻击时，人工神经网络和一些机器学习算法可以有效地为识别和避免遭到这些攻击提供帮助。此外，在物联网安全方面，也有研究者对机器学习和深度学习在这一领域的适用性上表现出高度的兴趣。

纵观当前的相关研究成果，在“人工智能影响下的密码学”这一议题中，

“人工智能嵌入密码学”阶段受到了大量的关注，尤其是如何在密码学中使用神经网络安全地加密和解密信息、如何在图像编码中使用神经网络、机器学习在密码学中的应用等（有关“人工智能影响下的密码学”详细解析请见寰球密码简报（第 134 期）| 密码与人工智能：从共同演化到量子革命）。不过，大部分的研究尚未涵盖量子计算对人工智能影响下密码学发展的影响。同时，一些研究中也指出了神经网络和量子密码学之间的一些不兼容性，这主要是由于基于神经网络的加密算法不能与任何基于海森堡原理和光子偏振的量子加解密技术兼容。

三、人工智能带来的威胁

得益于云计算、大数据和物联网（IoT）的发展，人工智能在近年来亦发

展迅速，并在包括机器学习、推荐引擎和自然语言处理等方面都得到了广泛的研究。在这样的发展态势之下，许多相关的技术开发也都取得不小的进展，如在自动驾驶汽车和家用无线设备的开发中就能看到不少人工智能发展的影子。此外，一些更为复杂的机械系统开发（如基于人工智能技术开发的脑机接口、人机交互设备等）也正随着人工智能的发展而被纷纷提上日程。

然而，在没有合适的加密技术辅助的情况下，人工智能可能会在带来便利的同时带来一系列的灾难。而解决安全问题的关键就是密码学，以及如何将密码学应用于人工智能。就当前情况而言，相关技术主要集中在机器学习、非对称加密、安全外包处理和多边信息处理等方面，可验证技术（保证人工智能系统完整性和精度的一种方式）的重要性也日益增强。但总体上，目前仍缺少能够为人工智能提供可实现安全性和准确性之间平衡的加密技术，已有的相关解决方案仍存在大量的计算或策略依赖，这对它们的实用性和适用性都产生了重大影响。

目前，人工智能（尤其是其中的深度学习技术）已经越来越多地被用于数据处理、语音识别和图像分析等应用场景，如计算机编程、网络安全、刑事司法、自动化等领域已经越来越依赖于人工智能技术。然而，在人工智能的快速发展和创新中，隐藏在其中的负面影响也逐渐显现，如何应对与人工智能相关的安全威胁和犯罪行为成为了需要更多关注的问题。

（一）人工智能数字取证

有关人工智能的犯罪可以被分为

在没有合适的加密技术辅助的情况下，人工智能可能会在带来便利的同时带来一系列灾难。而解决安全问题的关键就是密码学，以及如何将密码学应用于人工智能。

两类，即以人工智能作为工具犯罪和将人工智能作为目标的犯罪。其中，以人工智能作为工具的犯罪实际上是对已有犯罪行为的延展，如网络钓鱼、计算机黑客入侵等；将人工智能作为目标的犯罪则主要集中在破坏人工智能本身的功能，如用对抗性攻击令人工智能作出误判等。针对以上犯罪行为进行调查时，就可以进行针对人工智能的取证，其本质是对人工智能本身的研究，而不是单纯将人工智能技术用于调查，其与传统的网络证据获取有所区别。

值得注意的是，一些研究表明，拥有像人类一样行为的在线社交机器人（Social Bot）已成为研究人工智能取证的关键之一。虽然在线社交机器人最初的设计目的是提升人们之间的合作效率，但其已经越来越多地被用于不法目的，包括黑客攻击、网络盗窃和对在线社交媒体活动的渗透等，并因此在影响和煽动公众情绪方面造成不良影响。

（二）人工智能有害使用

当前，一些恶意攻击者已经开始尝试将人工智能武器化，并将其直接运用于犯罪行动中。例如，一些网络犯罪分子已经可以利用人工智能技术模仿受害者家人和同事的声音，或使用人工智能评估和识别系统中的漏洞以进行攻击，并以此进行诈骗或勒索。

在针对人工智能有害使用的研究中，一些研究人员关注到了这其中三个主要的变化：已有风险的增长、新风险的出现及已有风险特征的改变。当犯罪分子能够灵活而熟练地使用人工智能时，由于犯罪成本的下降，犯罪分子将更容易在短时间内实施更大规模的攻击（如大规模的网络钓鱼骗局），这将导致原本已有风险的扩散

和增长；随着人工智能的发展，原先无法实现的新型犯罪手段将越来越多（如语音模仿和多无人机控制等），这将导致新风险的出现；当人工智能攻击变得更加普遍时，原有的风险性质和特征将发生改变。

此外，也有一些研究将人工智能有害使用/犯罪以更独特的视角进行了分类，即分为对资本市场的侵蚀、商业行为犯罪（如市场投机、价格管制等）、有害或致命内容、对人的犯罪（如言语攻击等）、性犯罪、盗窃、欺诈等。这些研究人员在评估人工智能安全威胁时更专注于人的表现、义务、责任和心理，例如，人工智能带来的道德威胁意味着人工智能有可能影响人们的精神状态，从而促进或导致犯罪。

（三）人工智能隐私问题

除了对人工智能取证和人工智能有害使用的关注，也有一些研究集中在人工智能带来的隐私问题上。

其中一些研究强调，人工智能在医疗、金融和教育领域的应用可能会出现隐私问题。由于训练数据的数量和质量对人工智能的有效性有重大影响，人工智能开发人员的目标往往是“尽可能多地积累信息”，而这种对大量数据的获取行为是存在不可忽视的隐私和保密风险的。针对这一问题的解决方案，目前主要依赖于如《通用数据保护条例》（GDPR）等相关的法律和规定来进行规范，即当人工智能管理下的内容属于“个人信息”的范畴时，就可以被纳入相关立法的调整范围。

（四）人工智能安全威胁

除了人工智能在虚拟环境中带来的问题，也有一些研究人员开始思考其对在真实环境中的安全威胁，尤其是随着物联网的普及，一些人相信，

人工智能带来的安全问题已经超越了网络的限制，即罪犯可以通过控制人工智能系统对受害者进行直接的人身攻击。而同时，当前的法律和伦理体系下却很难确定谁在道德和法律上对这种伤害负有真正的责任。例如，一些研究表明，目前一些汽车电子控制系统中的缺陷，就可能存在其被恶意入侵、攻击和控制的风险。一系列的相关试验显示，如图像识别、漏洞扫描等技术都有可能降低人工智能系统的有效性，而鉴于人工智能目前在教育、无人驾驶、医疗等多领域上越来越广泛的应用，这些攻击可能最终会造成对受害者更直接的伤害。

另外，值得注意的是，也有一些相关业内人士和研究人员特别关注到了量子计算和人工智能的共同发展。他们认为，当前最广泛使用的加密方法（包括 Diffie-Hellman 和 RSA 等算法）都将因为量子计算对人工智能发展的加持而很快被击败，而应对这一挑战需要有足够强大的解决方案，并最终形成一个能够针对许多解决方案的标准化程序，以应对人工智能在各个层面上带来的威胁，可上述的一切目前都尚未出现。

四、结论

人工智能对于密码学的影响具有两面性。一方面，由于人工智能算法可以根据大量的数据进行学习，并将学习结果应用于识别加密算法中的脆弱性并对其实施攻击，这可能会大大减少破解加密所需的时间和资源，导致原本被加密保护的敏感数据安全性被破坏。另一方面，人工智能也可以用于增强密码学，例如，通过使用机器学习来检测攻击并相应地调整算法以抵抗攻击、利用人工智能技术开发

更安全的加密算法、检测加密数据中的异常情况，合理地使用人工智能将切实地提高加密的效率和效果。

为了降低人工智能在密码学中的风险，谨慎地思考人工智能对密码学的益处和潜在风险，在利用益处的同时制定策略来减轻风险是至关重要的。从技术层面来说，在使用人工智能来增强现有加密算法的同时，应同步开发新的算法来抵抗人工智能的攻击；从投资层面来说，应更多地注重对于相关研究和开发的资金投入，以更好地了解人工智能在密码学中的潜在风险，并制定策略来减轻这些风险。

当前，针对人工智能在破解加密中的应用所带来的威胁仍然需要展开更多的研究和关注。该篇文献综述在总结了目前相关学界的多种观点后，为人工智能领域的从业者提供了以下三个建议。

首先，即使在人工智能被合法使用的情况下，科学家和技术人员也应该意识到，由于人工智能的双重用途，其随

时可能被用于犯罪。因此，基于人工智能的两面性，涉及这一领域的主体需以较为严格的道德标准来进行筛选。

其次，由于人工智能可以完成人类（传统程序员）无法完成的事，人工智能和密码学领域的从业人员、研究人员与各个领域的专家之间有必要进行持续的研究和合作，以充分调查人工智能对密码学带来的影响，并开发更强大的安全解决方案来应对未来的挑战，最大限度地减少人工智能的潜在安全问题并应对人工智能犯罪带来的挑战，以确保人工智能的发展不会危及数字系统的安全。

最后，人工智能犯罪实际上与网络犯罪密切相关，而在网络犯罪中，具有双刃剑性质的信息通信技术既是犯罪的根源，也可以成为对抗犯罪的有力武器。也正因如此，相关安全部门应更多地考虑利用人工智能作为保障安全的工具，以此来应对已经出现的人工智能犯罪。（本文来源：《寰球密码法律政策发展动态简报》）

关于作者

洪璐 方婷

苏州信息安全法学所密码法研究中心

西安交通大学苏州信息安全法学研究所，由我国著名网络与信息安全法律专家马民虎教授领衔。法学所负责密码法治实践创新基地的具体建设，致力于国家信息安全政策和法律重大专题项目研究、国家部委信息安全政策法律咨询、企业信息安全合规研究等。

安全团队指南： 如何创建网络安全的“谷歌地图”

作者 约翰·莫雷洛

如何将网络安全从纸质地图和零碎分析转变为集成、步调一致、全面、可实时导航的系统？

网络安全不仅仅是防火墙和杀毒软件，还要了解您的防御系统、人员和流程是如何协同工作的。就像谷歌地图彻底改变了导航方式一样，流程映射也能彻底改变您了解和管理安全状况的方式。

过去，我们常常拿着纸质地图去新的地方。这既危险又不方便，边开车边看地图意味着两件事都做不好。后来，Garmin 和 TomTom 等公司推出了看似神奇的逐向导航 GPS 系统。这些系统售价数百美元，堪称奢侈品。谷歌地图问世后，纸质地图突然被废弃使用，而“逐向导航”（“turn-by-turn”）变得不可或缺。我们中的大多数人已经无法想象，在计划旅行、步行或驾车导航新城市时，如果没有网络地图，将如何生活。新的用例出现了，如今，我们使用网络地图就像使用搜索引擎一样，可以过滤评级、营业时间、食品或商品种类。

如今，大多数团队在网络安全领域都使用纸质地图。有些团队使用电子表格，并进行手动更新。还有一些团队使用的是可能包含一些自动化元素的仪表盘。在最糟糕的情况下，他们不得不解析日志文件，手动将一连串事件或迹象拼凑在一起，以绘制网络安全流程图。就像事件响应一样，周而复始，每一次事件分析过程都费时费力，令人苦不堪

言。这也是当今网络安全仍然低效、不精确的重要原因。

那么，我们如何才能将网络安全从纸质地图和零碎分析转变为一个集成、步调一致、全面、可实时导航的系统，从而为我们提供相当于逐向导航的可视性和规划呢？

采用流程制图思维，建立实时可视化的安全旅程——安全工作流程的谷歌地图。

第一步：定义关键路径

您不需要绘制所有内容，只需要选择谷歌地图工具绘制跟踪安全流程或工作流程的元素。

确定关键流程：从最重要的安全工作流程开始。这可能包括事件响应、漏洞管理、威胁猎取或合规性审计。

绘制地形图：将每个流程分解为各个步骤。谁参与其中？采取了哪些行动？使用了哪些工具？要细致入微。

第二步：绘制网络安全地图

谷歌地图的魅力在于其强大的可视化架构，对于网络安全流程而言，可以简化“选择你自己的冒险”。

选择制图工具：从具有动态节点的简单流程图软件到专门的网络安全流程制图平台，有很多选择。理想的工具可让您创建动态、交互式流程图，并可根

据任何关键属性（角色、条件、位置、流程类型）进行实时更新和过滤。

与工具集成：将地图链接到信息安全管理（SIEM）、工单系统、聊天工具、电子邮件和安全协调工具等。这样，您的地图就能反映安全操作的实时状态，并直观地显示谁做了什么及何时发生。集成应提供可视化流程中的交互时间轴，让您轻松快速地浏览感兴趣的流程。

绘制地图：将每个流程的步骤连接成一个可视化流程。使用彩色编码突出显示不同的团队、状态或潜在瓶颈。添加注释和注解以提供上下文。地图必须提供完整的活动链及嵌套动作和反应的可见性，以正确捕捉网络安全团队的工作导航方式。

第三步：使用地图优化安全流程

现在您已经有了关键安全流程的地图和可视化景观，您可以部署一个强大的安全商业智能（BI）工具，让您直观地检查和分析不同的过程如何导致不同的结果。这是安全优化最关键的部分——优化人为因素，关注人们的实际

工作（而不是仪表盘或证明上所说的工作）。

按类型分析流量模式：绘制特定类型事件在流程中的流程图，以及不同任务的执行方式。延误在哪里？是否有意想不到的迂回？人们是否跳过步骤或不遵守规定？还是他们自己优化了流程并提高了效率？

调查具体事件：使用您的地图调查特定事件或行动，从应对入侵迹象到修补高风险性零日漏洞。查看正在发生的事件和遗漏的事件。

识别流程风险并优化流程手册：更新您的流程，以简化工作流程、消除不必要的步骤，并自动执行重复性任务。使用地图测试和验证您的更改。

绘制永无止境的安全演进图

尽管谷歌地图是一款出色的产品，但我们都会在使用过程中遇到一些错误——商店或餐馆关门了，距离稍有偏差，方向指示告诉您在有“禁止左转”标志的十字路口左转。同样，您的安全流程图也需要不断发展，以跟上组织、工具和流程的变化。您的安全环境在不断变化，因此您需要定期审查和更新您的地图，以反映新的威胁、工具或流程。

安全本身是一个过程，而不是产品。人类处理视觉信息的效率要高于其他任何形式的信息。我们还进化成了寻路者，天生就有绘制地图的意识。在线制图工具成为历史上最受欢迎的应用程序之一是有原因的。通过将这些经验应用到网络安全中，我们可以利用地图的力量来提高网络安全团队的效率，并最终清楚地了解对于良好安全而言最重要的信息——意图和行动地图，让我们能够重温、学习和改进。安

关于作者

约翰·莫雷洛

Gutsy 公司联合创始人兼首席技术官。Gutsy 采用新的方式进行安全治理，首次将流程挖掘应用于网络防护，以帮助 CISO 将安全视为相互关联、带来结果的系统和事件，而不是单独的 Settings 和检测。



「零事故」安服团队重磅力作

900余次实战攻防演练的经验总结
从红队、蓝队、紫队视角全面解读攻防演练



扫码购买

奇安信连续三年位居
“中国网安产业竞争力50强”
第一名



6月20日，中国网络安全产业联盟（CCIA）
公布“2023年中国网安产业竞争力50强”榜单，
凭借扎实的技术实力和领先的市场表现，
奇安信连续三年高居榜单第一名。



“2022年中国网安产业竞争力50强”榜单

TOP15 公司名称

- 1 奇安信科技集团股份有限公司
- 2 深信服科技股份有限公司
- 3 启明星辰信息技术集团股份有限公司
- 4 华为技术有限公司
- 5 天融信科技集团股份有限公司
- 6 绿盟科技集团股份有限公司
- 7 腾讯云计算（北京）有限责任公司
- 8 新华三技术有限公司
- 9 阿里云计算有限公司
- 10 杭州安恒信息技术股份有限公司
- 11 三六零数字安全科技集团有限公司
- 12 亚信安全科技股份有限公司
- 13 中电科网络安全科技股份有限公司
- 14 杭州迪普科技股份有限公司
- 15 山石网科通信技术股份有限公司