

2025年

网络空间安全漏洞态势 分析研究报告



目录

第一章 前言.....	3
第二章 2025 年漏洞技术趋势	5
2.1 “智能编程” 安全漏洞不可避免.....	6
2.2 基于漏洞的供应链攻击潜伏性更强	7
2.3 利用大模型的漏洞挖掘效率大幅提升.....	9
2.4 大模型自身漏洞成为新的攻击面.....	10
2.5 关键基础设施漏洞攻击危害加剧.....	11
2.6 加密货币交易所漏洞攻击不断出现	12
第三章 国家漏洞库概况	13
3.1 漏洞威胁等级统计	14
3.2 漏洞影响对象类型统计	15
3.3 漏洞产生原因统计	17
3.4 漏洞引发威胁统计	18
3.5 行业漏洞收录统计	19
3.6 漏洞修复情况统计	20
3.7 漏洞增长趋势	21
第四章 在野利用漏洞概况.....	22
4.1 漏洞影响厂商分布情况.....	23

4.2	漏洞影响平台产品分类.....	24
4.3	漏洞类型统计概况.....	25
4.4	EXP 公开情况统计.....	29
第五章	最值得关注 TOP10 高危漏洞情况	31
5.1	年度 TOP10 高危漏洞	31
5.2	年度漏洞预警 TOP10 漏洞回顾	33
第六章	2026 年漏洞技术走向	40
6.1	“智能编程” 安全漏洞不可避免.....	41
6.2	供应链的漏洞攻击仍威胁严重	42
6.3	AI 漏洞挖掘与利用效率将进一步提升	43
6.4	大模型自身的漏洞将成为新的攻击热点	44
第七章	总结.....	46
7.1	安全防护建议.....	47

第一章 前言

2025 年，我们看到网络安全领域出现了一些关键变化。人工智能等技术正被更深入地融合进攻击手段中，漏洞利用的速度和规模进一步加快，甚至在漏洞公开披露的极短时间内就被恶意使用。防御的一方常常感到时间不够用，例如从漏洞公开到遭受攻击，有时窗口期被压缩到几乎以小时计算。漏洞的类型也在悄然集中，例如无需认证且高影响力的漏洞。反序列化漏洞依然是许多攻击中的利器，内核与本地提权漏洞也频繁出现在攻击链条里。值得注意的是，随着 AI 组件的广泛应用，与之相关的安全漏洞明显增多，大模型开发框架中的风险点也逐渐暴露。智能体技术在推动效率提升的同时，也正在改变攻防的形态，自动化对抗正在变得普遍。另一方面，物联网设备的安全问题依然突出，供应链环节的薄弱和关键基础设施面临的勒索威胁，都让漏洞不再只是技术层面的问题，更成为安全实践中持续博弈的焦点。在这些趋势下，国内也在持续推动“两高一弱”等重点安全治理工作，加强跨部门协作，尝试在安全与业务发展之间寻找更可持续的平衡。

2025 年全球频发重大网络安全事件，涵盖 AI、供应链、加密货币、Web 开发、游戏及短视频领域等，凸显多领域安全防线均存短板，攻击危害愈发严重。1 月 AI 平台 DeepSeek 上线一周即遭大规模攻击；2 月 Bybit 被盗走超千亿加密资产；3 月 GitHub 热门 Actions 遭供应链入侵致敏感信息泄露；12 月 React 核心组件曝高危远程执行漏洞、育碧因数据库漏洞泄露海量源码、某短视频平台遇规模化黑灰产攻击致市值大跌，既说明新兴技术落地安全适配不足，也体现传统网络防护漏洞仍为高危隐患，同时针对性攻击、供应链渗透已成黑客主流手段，各行业面临的网络安全挑战愈发严峻且无死角。

作为中国领先的网络安全、大数据与云服务提供商，天融信始终以捍卫国家网络空间安全为己任，创新超越，致力于成为民族安全产业的领导者、领先安全技术的创造者和数字时代安全的赋能者。为了更好地了解网络空间安全漏洞的发展趋势，并采取适当的措施

应对漏洞威胁，特发布《2025 年网络空间安全漏洞态势分析研究报告》。本报告旨在全面梳理和分析 2025 年网络空间安全漏洞的态势和特点，并结合实例剖析典型安全事件的成因和影响，预测 2026 年漏洞技术走向以期为政府、企业和个人提供有针对性的防护建议和措施建议共同应对网络安全挑战维护网络空间的安全与稳定。

本报告重点内容共分四个部分，第一部分为 2025 年漏洞技术趋势，通过对智能编程、供应链安全、AI 大模型应用与自身安全、以及关键基础设施和交易所攻击案例的综合分析，指出智能编程普及导致代码控制力下降、供应链攻击潜伏性增强、大模型成为新攻击面并显著提升攻防效率、以及针对关基和交易所的攻击呈现出利用底层协议与供应链漏洞控制物理世界或实施高度隐蔽欺诈的新特点。

第二部分为 2025 年度漏洞趋势，通过对国家漏洞库及前 100 个在野利用漏洞进行综合分析而产生。据 CNVD 公开数据显示，2024 年共披露漏洞 18782 个，2025 年共披露漏洞 19576 个，同比增加 4.23%。这可能表明新的攻击面被发现。其中，低危漏洞占 5.65%，中危漏洞占 45.47%，高危漏洞占 48.88%。结合 AI 武器化，漏洞利用“零日化”趋势，漏洞治理已从技术层面升级为关乎企业生存的战略要务。

第三部分为年度最值得关注的十大高危漏洞，天融信阿尔法实验室持续监测全球漏洞态势，累计捕获并分析漏洞情报数万条。通过对海量信息的实时追踪与快速研判，实验室成功预警并协助处置了多起突发性的极高危漏洞安全事件。基于漏洞的影响范围、潜在危害程度及实际威胁等级，系统性地梳理并发布了年度最具代表性的十大高危漏洞。本年度重点漏洞涵盖 React 服务端组件、Microsoft 产品、Redis、Sudo、Kubernetes、Langflow、苹果设备 WhatsApp 及 WinRAR 等多类载体，涉及远程代码执行、权限提升、反序列化、认证绕过、路径穿越等多种漏洞类型，均对系统安全、数据安全构成严重威胁。在每一起极高危漏洞爆发后，天融信阿尔法实验室均第一时间启动应急响应流程，迅

速完成漏洞复现、影响范围分析与攻击链研判。通过发布及时、详尽的风险通告，为客户提供有效的临时缓解措施与长期修复方案。

第四部分为 2026 年漏洞技术走向，通过对智能编程、供应链安全、AI 大模型应用与自身安全等方面的综合研判，预测全年漏洞总量将保持高位甚至小幅增长，指出智能编程模式的持续迭代将带来代码安全性挑战，供应链攻击仍将是高回报的主要威胁手段，AI 将深度融入攻防两端显著提升漏洞挖掘与修复效率，同时大模型及其基础设施由于安全体系尚不完善，将成为新的攻击热点和研究重点。

企业安全部门应该提升网络安全威胁感知能力，建立有效的监测预警体系，并制定应急指挥计划，加强对威胁的发现、监测、预警和应对能力。以便在网络攻击发生时快速作出应对。此外，还需要强化攻击溯源能力，重点建设辅助攻击溯源相关能力（如授权控制、流量记录），使得对网络攻击的溯源工作更加轻松，以便有效地防御和应对来自外部的网络攻击。

天融信阿尔法实验室秉承攻防一体的理念，以保卫国家网络空间安全为己任，在未来的工作中将持续针对网络空间漏洞进行实时侦测，并灵活应对和防护突发漏洞的产生，攻防相结合，为国家及客户安全进行全方位赋能。

第二章 2025 年漏洞技术趋势

2025 年在漏洞相关安全技术方面，存在以下特点。

（1）智能编程带来了工作效率提升，也造成了代码控制力下降。在此背景下，当智能编程模式的日活率进一步提升，代码控制力下降叠加大量的 AI 生成代码，整体的安全漏洞态势仍不容乐观；

（2）供应链攻击叠加开源代码使得潜伏性更强，更具针对性；

(3) AI 大模型自身安全成为新的攻击面，值得安全研究员深度探索；

(4) AI 大模型使得攻防速度大幅提升，利用大模型进行安全测试、漏洞挖掘、漏洞利用、告警降噪、流量识别等，已被证明都具有良好应用；

(5) 网络攻击甚至是网络战争有了新的特点，针对关基、交易所等重要目标的攻击层出不穷。

2.1 “智能编程”安全漏洞不可避免

在软件开发领域，AI 大模型正推动智能编程进入一个全新的时代。回顾 2024 年，许多开发者已开始借助大模型的代码生成能力，将其生成的函数级代码手动整合到项目中，实现初步的辅助编程。这一阶段的 AI 编程，仍停留在小规模、片段化的代码生成与应用。

进入 2025 年，随着“深度思考”机制的显著进步，AI 大模型的推理能力大幅增强，推动更智能的 Vibe Coding 模式走向现实。在这一模式下，开发者无需逐行编写代码或强记语法，而是直接通过自然语言（如中文、英文）向 AI 发出指令，依托 VS Code、Cursor 等集成开发环境，由大模型自动完成编码、测试，并输出可直接运行的软件产品。借助更强大的上下文理解与复杂逻辑处理能力，Vibe Coding 已能应对真实软件工程需求，而不仅仅是编写演示原型。这种模式的演进，显著提升了开发效率，以往需数小时完成的任务，如今可在几分钟内解决。

Stack Overflow 于 2025 年 11 月发布的《2025 年度开发者调查报告》显示，共有来自 177 个国家的超 49,000 名开发者参与，涵盖 314 项技术相关议题。调查指出，2025 年有 84% 的受访者正在使用人工智能工具。

GitHub Copilot 是由 GitHub 与 OpenAI 共同开发的人工智能编程助手，能够在编程时提供实时代码补全与建议，帮助开发者提升编码效率。它通过分析海量代码库学习编

程模式，并已集成到 VS Code、Visual Studio、JetBrains IDEs 等主流编辑器中，能够依据上下文智能生成函数、代码块乃至测试用例。

微软作为 GitHub 的母公司，会在财报中披露其核心产品数据。2025 年中期数据显示，GitHub Copilot 用户数（含试用与付费）已突破 2000 万，且超过 90% 的财富 100 强企业已使用该产品。作为背景参考，据 Evans Data / Statista 2024 年统计，全球专业软件开发者数量约为 2870 万人。

在此背景下，当智能编程模式的日活率进一步提升，代码控制力下降叠加大量的 AI 生成代码，软件产品的安全漏洞态势不容乐观。

2.2 基于漏洞的供应链攻击潜伏性更强

2025 年底披露的 React2Shell (CVE-2025-55182)漏洞案例中，React 组件的 React Server Components (RSC)功能，允许 React 组件仅在服务端运行。漏洞根本原因分析结果表明，RSC 功能依靠 Flight 协议实现服务端与客户端的无缝交互，由于 Flight 协议的过滤不完善，导致 RSC 功能底层实现存在反序列化漏洞，能够远程执行任意代码。该漏洞利用难度低、无需权限、可执行任意代码，最重要的是这个漏洞几乎影响极为广泛，几乎所有依赖 React19 及其相关版本（主要是集成了 RSC 功能）的 Web 系统均受影响。RSC 是 React 团队提出并定义的一套技术规范和底层能力，而 Next.js 是第一个将其完美封装并让开发者能直接用于生产环境的框架。由于 RSC 存在这个漏洞，使得 Next.js 受到间接影响。由于 React、Next.js 都是非常流行的 Web 应用基础设施，这也导致大量 Web 应用受到影响。

开源不等于安全。多数人误以为开源等于安全，实际上这是严重误解。形成这个误解源于，互联网上存在一个观点：“只要有足够多的关注，所有的 BUG 都将无处遁形”。意

思是说，开源软件因为代码是公开的，全世界的开发者和安全专家都能看到它，所以漏洞会被迅速发现并修补。但实际上并不完全是这样。考虑以下三个问题。

- 第一个问题是大部分开源代码，关注的人并不多。只有像 Linux 内核、Nginx、Apache HTTPD 这类明星项目才有成千上万的安全研究员审计。许多被广泛使用的开源库，可能只有几个维护者在业余时间维护；
- 第二个问题是，既然代码是公开的，不仅防御者可以通过审计代码来发现并修复漏洞，攻击者也可以通过审计来寻找缺陷，甚至可能比防御者更快的发现 Oday 漏洞。值得注意的是在闭源软件中，为进行安全审计，攻击者还必须进行逆向分析。而在开源系统中，他们可以直接审计源码寻找弱点。这其实对攻击者更有利；
- 第三个是供应链安全问题。现代软件开发类似于搭积木，90%的代码可能来自第三方开源库，剩下约 10%的业务代码由开发者编写。如果攻击者控制了底层开源库，所有使用这个库的软件都会受到影响。2024 年的 XZ Utils 后门事件就是典型。前文提到的 React2Shell，虽然不是典型的“投毒型”供应链攻击，但绝对属于特殊的供应链安全问题。

所以说，“开源软件就是安全的”绝对是一个误解。React2Shell 是一个典型的供应链安全灾难，虽然不是人为投毒，但它暴露了现代软件开发对“巨型框架”过度依赖的脆弱性。开源软件并不因为“用的人多”就天然安全。React 这种由顶级大厂（Meta/Vercel）维护，全世界都在用的项目，依然可能出现满分的核弹漏洞，使得下游 Web 系统受其影响。

此类事件的披露，对安全开发和安全运维工作而言，具有重要启示作用。对于软件开发，应当警惕供应链安全问题，践行最小化原则。也就是说，为了杜绝供应链安全问题，应严格限制组件数量，从物理上切断攻击路径。没有代码，就没有漏洞，也就完成了攻击面收敛。对于安全运维，同样是攻击面收敛，减少业务运行环境中不必要的组件依赖。同

时应该规范漏洞修复工作，有效梳理软件产品及运行环境的成分，分级分类及时更新重要安全漏洞。

2.3 利用大模型的漏洞挖掘效率大幅提升

在漏洞挖掘与利用方面，AI 大模型的应用也已经非常普遍。根据 Bugcrowd 发布的《2025 年度黑客洞察报告》，2025 年在漏洞挖掘过程中使用过 AI 工具的安全研究员比例，已从 2024 年的 77% 上升至 87%，这意味着使用 AI 的研究人员已从“大部分”变为“绝大多数”。

报告同时指出，安全研究人员使用 AI 大模型的用途，主要集中在以下三个方面：

- 1. 编写辅助脚本：**这是最常见的用途。研究人员可通过自然语言提出需求，例如：“帮我写一个 Shell 脚本，遍历文件夹中的所有文件，提取文件中的 Token 信息。”此前编写和调试类似脚本可能需要 30 分钟，而现在 AI 可在 1 分钟内完成，极大提升了自动化效率。
- 2. 代码解释与分析：**AI 在此扮演“代码助手”角色，帮助研究人员快速理解代码逻辑。例如，用户可输入：“解释一下这段 Java 代码在干什么”，从而迅速把握代码功能与潜在风险。
- 3. 撰写漏洞报告：**研究人员可指示 AI 将漏洞发现过程转化为结构清晰、语言专业的报告，例如：“把这个命令执行的过程，整理成一份包含关键步骤与修复建议的漏洞报告。”借助 AI 的语言生成能力，报告不仅质量更高，还可按需生成多语言版本。

这些应用不仅显著提升了研究效率，也反映出 AI 正深度融入安全工作的全流程。

2.4 大模型自身漏洞成为新的攻击面

随着 AI 大模型的快速发展，其基础设施也呈现“野蛮生长”态势，尤其是模型的 API 的快速部署与集成过程中，暴露出大量新型安全隐患。这使得 AI 应用自身正成为一个新的、复合型的攻击面。一方面，上层的大模型存在特有的安全问题，如提示词注入、训练数据泄露等；另一方面，底层的支撑基础设施中，传统类型的安全漏洞（如配置错误、权限缺陷、未授权访问等）也呈现爆炸式增长。

以 2024-2025 年间企业广泛部署的 RAG（检索增强生成）系统为例。这类系统依赖向量数据库存储和处理知识库，然而大量开发团队为测试方便，未设置严格鉴权机制便直接将数据库端口暴露在公网。这意味着攻击者无需攻陷 Web 前端，仅通过网络扫描即可直接访问并窃取企业核心知识库，即所谓“拖库”风险。根据 ISC2 2025 年安全大会披露的信息，此类向量数据库暴露问题已成为 2025 年网络层漏洞报告的主要来源之一。以下是一组数据：

- 2025 年国内首次 AI 大模型实网众测活动结果显示，活动共对 10 家 AI 厂商的 15 款大模型及应用产品进行了测试。测试产品中既有基础大模型产品，也有垂域大模型产品，还有智能体、模型开发平台等相关应用产品，其中既包含单模态大模型，也涵盖多模态大模型，具有较广泛的代表性。测试累计发现各类安全漏洞 281 个，其中大模型特有漏洞 177 个，占比超过 60%。数据表明，当前 AI 大模型产品面临着大量传统安全领域之外的新兴安全风险。
- HackerOne 发布的第 9 届年度 Hacker-Powered 安全报告提供的数据中提到，2025 年 AI 漏洞报告同比增长 210%，其中提示注入漏洞同比增长 540%。
- Orca Security《2025 年云安全状态报告》显示，84%的组织在云端使用 AI 相关工具，其中 62%的组织环境中至少有一个易受攻击的 AI 包。

综合以上情况显示，当前 AI 大模型及其组件漏洞的数量在整体中的占比，较上一年有明显上升。

2.5 关键基础设施漏洞攻击危害加剧

2025 年，网络攻击不仅在数量上激增，更重要的是在质的层面发生了根本性的转变。利用安全漏洞针对关键基础设施甚至军事设施的攻击层出不穷。攻击者的耐心变强了。他们不再是砸碎玻璃抢东西的劫匪，反而是变成了拿着万能钥匙，静静坐在对方客厅里喝茶的隐形人，他们的目标或是将系统内的财产洗劫一空，或是在关键时刻给与致命一击。他们所利用的漏洞，也更加巧妙更加底层，甚至是基于漏洞链的组合、攻击目标的组合、攻击手段的组合，以实现战略目标。

近期在针对委内瑞拉的“绝对决心”军事行动中，网络战不再是辅助干扰，而是作为首轮打击力量，精准配合军事突袭。攻击展现了极高的同步性，造成委内瑞拉国家级断网，达成“物理瘫痪”与“信息封锁”的双重目标。具体情况从三方面展开。首先攻击者利用“信任模型”的逻辑摧毁针对加拉加斯电力网，攻击者利用工业控制协议缺乏零信任验证的设计缺陷。攻击者通过前期渗透进入内网后，向 SCADA 系统发送了符合协议规范的“合法关断指令”，直接导致物理断路器跳闸，造成全城停电；同时，针对通信基础设施，攻击者利用边界网关协议 (BGP) 的信任漏洞，通过控制关键路由节点，向全球互联网广播了针对委内瑞拉 IP 段的虚假路由信息。这不仅导致了 DDoS 级别的拥塞，更从逻辑上将委内瑞拉的国家网络剥离出国际互联网，切断了该国对外的所有通信链路；此外，针对委内瑞拉的防空雷达系统，有情报显示攻击者利用了硬件供应链中的预置后门。这些深埋在芯片或固件中的恶意逻辑，可能通过特定的无线电频率 (RF) 信号被外部激活，导致雷达系统在遭受空袭的关键窗口期发生“死锁”或重启。

此外，2025 年 11 月北欧某国的一家区域性能源聚合商披露一起“幽灵电网”事件。该机构负责管理分布在三个城市的分布式储能系统（电池阵列）。其远程管理网络使用了一套基于 Erlang 开发的定制化 SCADA 网关。攻击者扫描了公网 IP 段，并没有寻找常规的 Web 漏洞，而是专门通过 Banner 抓取，定位运行 Erlang 版本的 SSH 服务端口。攻击者利用 CVE-2025-32433 漏洞利用脚本，向 12 台边缘网关发送了特制数据包。仅耗时 3 秒，全部网关沦陷，获得 root 级 Shell 权限。最后攻击者下达指令，篡改了电池阵列的“过充保护阈值”，并强制开启全功率充电模式。事故造成严重后果。首先是物理损坏，由于保护机制失效，两处储能站的电池组发生严重过热，引发物理火灾和爆炸；其次电网波动：为了隔离故障，区域电网被迫紧急切断负荷，导致约 4 万户居民停电 6 小时；另外事后取证极其困难，因为恶意代码仅存在于内存中，设备重启后即消失（Fileless Malware），最初调查人员误认为是电池硬件故障，直到数周后在网络流量回溯中发现了异常的 SSH 握手包。

现在，网络攻击甚至是网络战争有了新的特点。攻击者的重心从“窃取数据”彻底转向了“利用底层协议和供应链漏洞来控制物理世界”。基础设施的安全性严重依赖其供应链的纯净度与底层协议的抗攻击性。

2.6 加密货币交易所漏洞攻击不断出现

参考 2025 年发生的 Bybit 交易所盗窃案等实际案例，显示出针对虚拟加密货币交易所的攻击已经发生了根本性的变化。攻击者避开坚固的区块链金库，利用安全漏洞猛攻脆弱的 Web2 基础设施与人员。

2025 年 2 月，Bybit 交易所遭 Lazarus 组织攻击，损失 15 亿美元。不同于传统的私钥窃取，此次攻击者实施了一场精心策划的定向社会工程学攻击，污染了

Safe{Wallet} 前端依赖的一个 UI 渲染库，并且使用了极其隐蔽的前端逻辑漏洞。

攻击者伪装成知名科技公司的招聘人员，长期接触并获取核心开发者的信任，随后以“高薪职位的代码测试”为诱饵，诱导开发者下载并运行含有恶意脚本的项目，隐蔽的恶意代码便瞬间窃取其 SSH 私钥与 GitHub 登录凭证，使黑客能够“夺舍”开发者身份，神不知鬼不觉地向代码仓库提交了被篡改的 UI 渲染逻辑，从而在源头完成了供应链投毒。恶意代码利用 JavaScript 的动态特性，Hook（挂钩）了交易生成函数。当管理员发起转账时，恶意脚本在内存中拦截了交易对象。它一方面向底层签名模块发送黑客的地址，另一方面强制控制 DOM 渲染层，在屏幕上“硬编码”显示合法的内部冷钱包地址。当 Bybit 的内部管理员在网页端进行多重签名操作时，被篡改的 UI 在前端欺骗了管理员。管理员眼睛看到的是“转账 10 ETH 到热钱包”，但底层的签名请求（Transaction Hash）实际上是“转账全部 ETH 到黑客地址”。

整个事件可以看到，攻击者不再死磕经过千锤百炼的冷热钱包私钥，因为很难突破，而是将战场前移到了代码构建阶段，攻击用户和管理员看到的界面。这种攻击具有高潜伏性。恶意代码在开发阶段就混入了，这导致交易所发布的官方版本软件本身就带有后门。

第三章 国家漏洞库概况

漏洞的统计与评判是评估网络安全情况的一个重要指标，天融信阿尔法实验室参考国家漏洞数据库数据,对 2025 年披露的漏洞进行了全方位的统计分析,下图是近十年漏洞数量走势图，从这个数据中可以看出，近十年来，CNVD 披露的漏洞数量总体呈现上升趋势。尤其是在 2018 年至 2021 年之间，漏洞数量大幅增加。2025 年来看，漏洞数量出现了上升，AI 驱动的自动化挖掘让漏洞暴露速度倍增，大模型支撑框架、物联网设备等新兴载体催生了新的漏洞类型，而供应链攻击的蔓延则让传统软件组件漏洞风险持续放大。更关键

的是，漏洞威胁的核心已从“数量”转向“质量与时效”。“披露即利用”成为新常态。

这一变化背后，是攻击者借助生成式 AI 降低技术门槛，实现攻击链工业化运作的冲击，

也是企业在 API 安全、供应链防护、AI 治理等新战场上面临的复合型挑战。

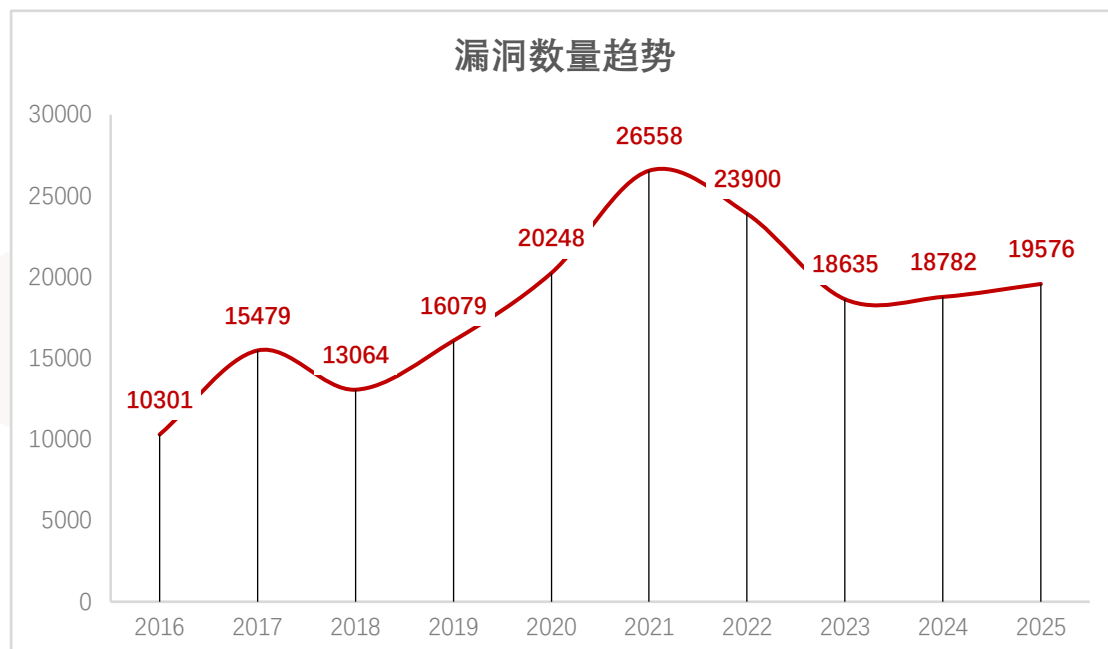


图 1 近十年漏洞数量走势图(数据来自于 CNVD 2026.1)

3.1 漏洞威胁等级统计

根据 2025 年 1-12 月漏洞引发威胁严重程度统计，其中低危漏洞占 5.65%，中危漏洞占 45.47%，高危漏洞占 48.88%。

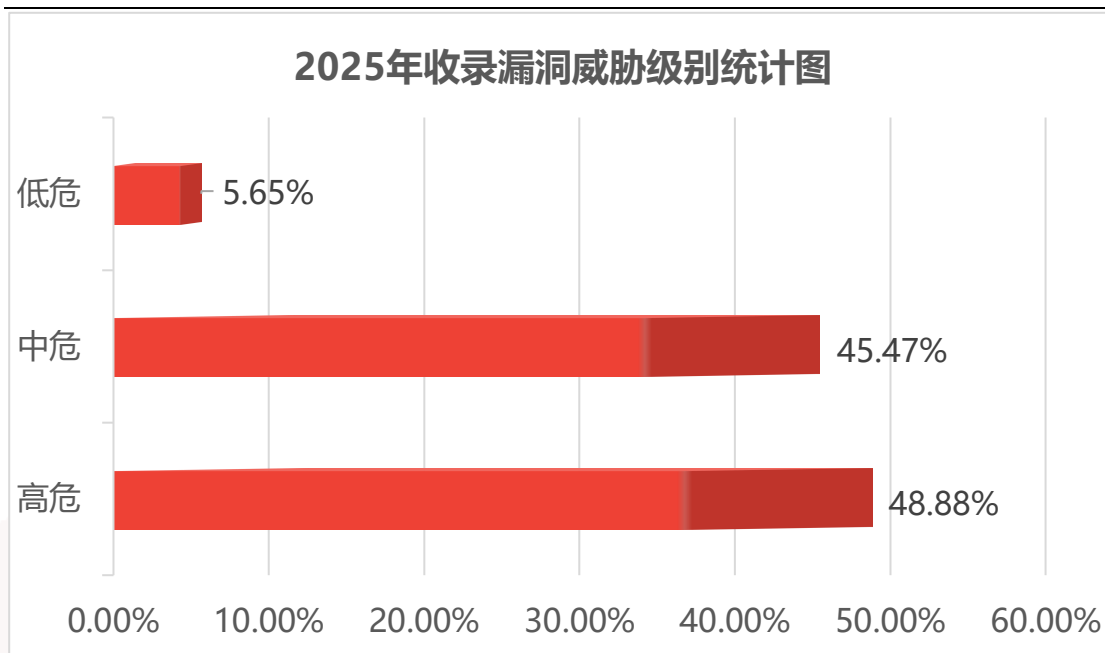


图2 2025年收录漏洞按威胁级别统计(数据来自于CNVD 2026.1)

分析 2024 年到 2025 年漏洞威胁统计的变化,低危漏洞的比例从上年的 4.20%上升到了 5.65%,资产暴露面扩大与基础配置疏漏增多。中危漏洞的比例则从 48.86%降至 45.47%,基础安全加固与合规治理初见成效。至于高危漏洞,则从 46.94%上升至 48.88%,漏洞武器化加速,供应链与基础组件漏洞成重灾区。

3.2 漏洞影响对象类型统计

根据 2025 年 1-12 月漏洞引发威胁统计,受影响的对象大致可分为九类:分别是 WEB 应用、应用程序、网络设备、操作系统、智能设备、工业控制系统、数据库、安全产品、车联网系统。其中 WEB 应用漏洞 38.70%,应用程序漏洞 27.90%,网络设备漏洞 19.10%,操作系统漏洞 5.50%,智能设备漏洞 4.20%,工业控制系统漏洞 2.80%,数据库系统漏洞 1.10%,安全产品漏洞 0.80%,车联网系统漏洞 0.01%。

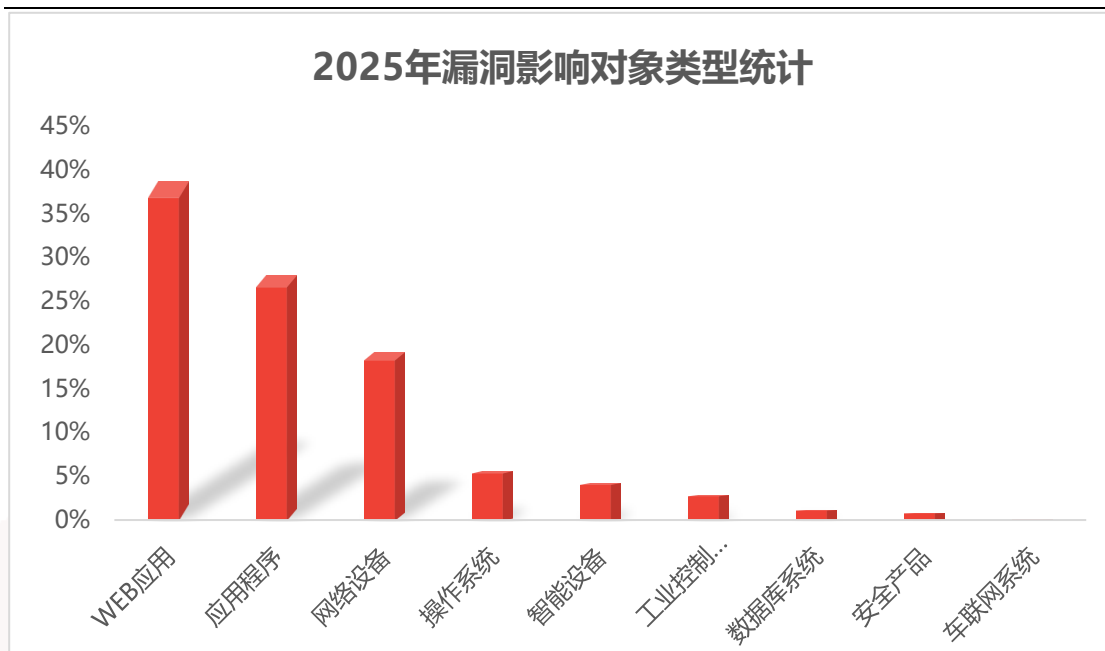


图3 2025年漏洞影响对象类型统计(数据来自于CNVD 2026.1)

2025年漏洞威胁统计显示，攻击重心正从传统应用层向底层基础设施与关键业务系统快速迁移，呈现出“攻防升级、焦点转移”的鲜明特点。

核心变化在于网络基础设施威胁激增。网络设备漏洞占比从13.80%大幅上升至19.10%，增幅达5.3个百分点，印证其已成为攻击者渗透内网的首要跳板。同时，承载关键业务的应用程序漏洞占比也从24.50%上升至27.90%。工业物联网成为风险增长最快的领域。工业控制系统漏洞占比从1.90%显著上升至2.80%，这与针对关键基础设施的勒索攻击常态化趋势直接对应，表明攻击意图正向物理破坏延伸。传统攻击面占比下降但威胁性质演变。WEB应用漏洞占比从48.10%下降至38.70%，但结合API成为主要泄露渠道的背景，其风险更集中、更致命。智能设备漏洞占比从5.00%微降至4.20%，反映治理初效与存量风险并存。其他方面，操作系统漏洞从4.00%至5.50%略有上升，而数据库、安全产品及车联网系统漏洞占比均出现小幅下降。

以上数据表明，防御方在应用层的加固迫使攻击者将火力更多投向网络基础架构和工业控制系统，网络安全的主战场正在向更底层、更关键的区域纵深拓展。

此外，车联网等新兴领域占比虽小，但作为“移动的智能设备”，其系统性风险不容忽视。数字化持续地赋能汽车产业，汽车已经从传统的出行工具逐渐演变成了新一代的移动数据中心和互联网服务创新的重要平台，智能网联汽车已成为汽车产业转型升级、交通方式变革、智慧城市建设的重要方向。未来，天融信将持续深耕车联网安全领域，全力构建智能网联汽车产业安全生态体系，为建设交通强国贡献企业力量。

3.3 漏洞产生原因统计

根据 2025 年 1-12 月漏洞产生原因的统计，设计错误导致的漏洞占比 56.10%，屈居首位，紧跟其后的是输入验证错误导致的漏洞占比 30.70%，位居第二，接着是边界条件错误导致的漏洞占比 10.60%，位居第三。后面的访问验证错误、其他错误、竞争条件和配置错误分别占比 2.10%、0.40%、0.10%、0.01%。

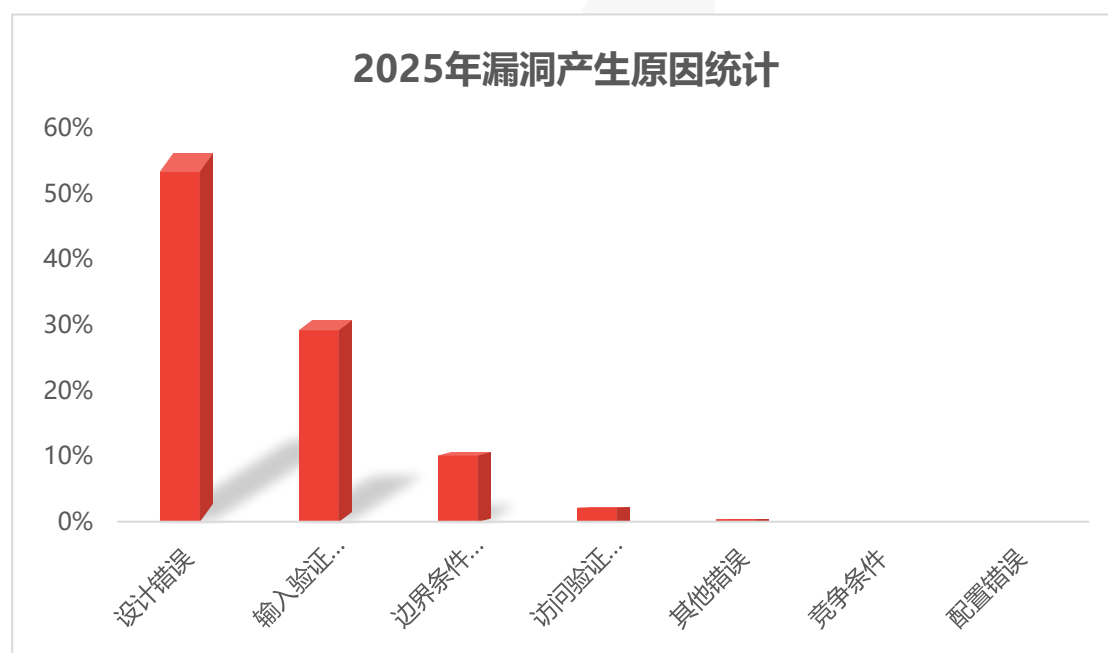


图 4 2025 年漏洞产生原因统计（数据来自于 CNVD 2026.1）

从漏洞产生原因的分布变化来看，设计错误导致的漏洞占比从 67.70% 降至 56.10%，虽仍位居首位，但占比回落反映出企业在产品研发阶段的安全设计意识有所提升；输入验

证错误占比从 26.20% 升至 30.70%，则体现出攻击者针对数据输入环节的渗透手段更趋多样，相关防护短板进一步凸显；边界条件错误占比从 3.50% 大幅攀升至 10.60%，说明复杂业务场景下的逻辑边界校验仍是安全研发的薄弱环节。此外，访问验证错误占比小幅上升，其他错误占比有所下降，竞争条件、配置错误占比保持稳定，整体态势映射出漏洞成因的结构性调整与攻防对抗的焦点转移。

3.4 漏洞引发威胁统计

根据 2025 年 1-12 月漏洞引发威胁统计，未授权的信息泄露占比 47.80% 位居首位，管理员访问权限获取占比 28.00% 位居第二，拒绝服务占比 12.60% 位居第三，后面的未授权的信息修改、普通用户访问权限获取、其它。占比分别是 11.50%、0.10%、0.10%。



图 5 2024 年漏洞引发威胁统计（数据来自于 CNVD 2026.1）

2025 年漏洞引发的威胁中，未授权信息泄露从 53.70% 降至 47.80% 仍居首位，占比下降显示防护有成效但仍是核心风险；管理员访问权限获取占比从 27.50% 到 28.00% 小幅上升，体现权限管控仍是攻防焦点；拒绝服务占比从 8.30% 到 12.60%、未授权信息

修改占比从 8.30%到 11.50%均显著攀升，与 AI 降攻击门槛、IoT 漏洞被滥用于 DDoS 等趋势呼应，而普通用户权限获取及其他占比均为 0.10% 无变化。结合 AI 落地加速、聊天机器人与智能体安全风险凸显、IoT 与供应链安全挑战加剧等态势，2026 年企业需强化数据防泄露、权限精细化管理、DDoS 防护与漏洞全生命周期治理，同时适配 AI 交互与爬虫验证等新安全场景，应对攻击目标转移与交互模式变革带来的新威胁。

在数据流转监测方面，天融信数据安全风险分析系统可以通过监控分析绘制全景图，并根据预设规则与算法，快速识别异常，持续分析数据流向，洞察潜在威胁，实现敏感数据监测和溯源，确保企业数据的安全合规使用。而天融信 API 安全审计系统则可以通过 API 审计、访问控制与权限审核、API 传输内容智能识别、数据流转监控与防护及溯源机制等模块，助力企业掌握自身数据资产分布以及数据流向路径监测，及时发现和处置潜在的数据安全风险。

3.5 行业漏洞收录统计

根据 2025 年 1-12 月行业漏洞统计，2025 年一共收录了电信行业漏洞 2328 个，移动互联网行业漏洞 521 个，工控行业漏洞 469 个。

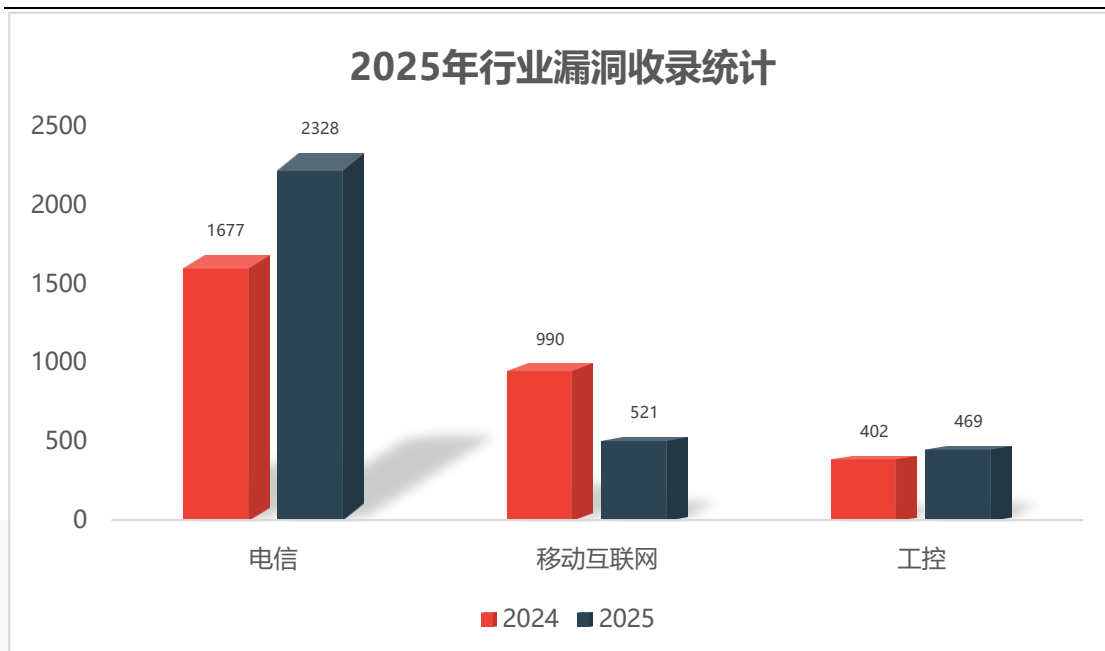


图 6 2025 年行业漏洞收录统计（数据来自 CNVD 2026.1）

2025 年的漏洞态势清晰表明，攻击者的战略重心正转向承载社会运转的核心命脉。电信行业因其基础设施的枢纽地位，已成为国家级网络对抗与战略渗透的关键跳板；移动互联网的攻击焦点则深入至碎片化的应用生态与供应链，利用 AI 实施更隐蔽的数据窃取与诈骗，此外为了实现 AI 自动化操作手机也可能会带来全新的攻击面；而工控系统漏洞的持续增长，直接映射出能源、制造等关键领域面临的、从网络勒索升级为物理破坏的极端现实威胁。

3.6 漏洞修复情况统计

根据 2025 年 1-12 月漏洞修复情况统计，漏洞数量排名前列的厂商与其修复数量如下图，其中 Apple、Oracle、Adobe、IBM、Cisco 漏洞修复率均为 100%，其余厂商中 Linux、Intel、Microsoft、Google、WordPress、吉翁电子、腾达，漏洞修复率分别为 99.85%、99.63%、99.42%、96.67%、71.47%、13.09%、7.53%。

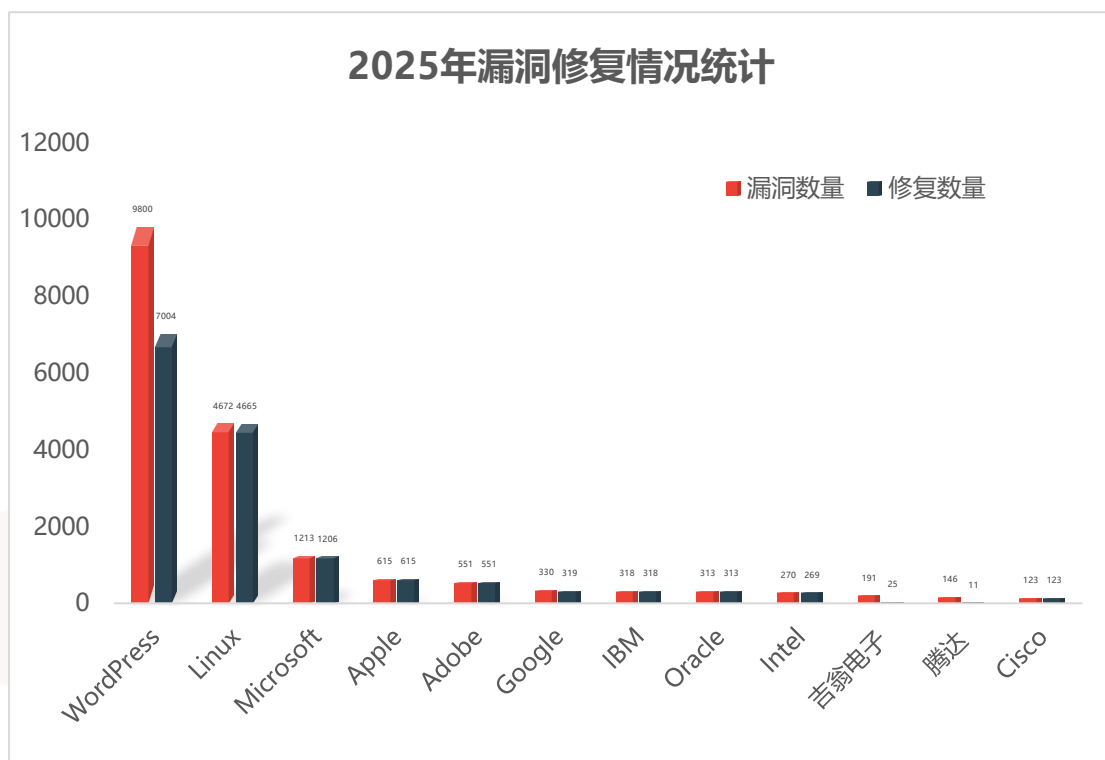


图 7 2025 年漏洞修复情况统计（数据来自 CNVD 2026.1）

数据显示，高修复率（如 Apple、Cisco 达 100%）展现了主流厂商成熟的漏洞响应能力，但修复的“速度”与“覆盖率”同等关键。与此同时，部分厂商（如腾达、吉翁电子）修复率显著偏低，这凸显了在广泛部署的硬件、IoT 及边缘设备领域，漏洞管理仍是严重短板。这些短板恰好是攻击者最易突破的入口，防御方必须将资源聚焦于保护承载社会运转的核心基础设施与供应链。

3.7 漏洞增长趋势

通过 CNVD 漏洞信息库对 2024、2025 漏洞公开数据显示，2024 年一共披露漏洞 18782 个，2025 年一共披露漏洞 19576 个，同比 2024 年增加 7.55%；2024 年高危漏洞 8499 个，2025 年高危漏洞 9569 个，同比 2024 年增加 12.59%；2024 年中危漏洞 8947 个，2025 年中危漏洞 8901 个，同比 2024 年减少 0.51%；2024 年低危漏洞 756 个，2025 年低危漏洞 1106 个，同比 2024 年增加 46.30%。

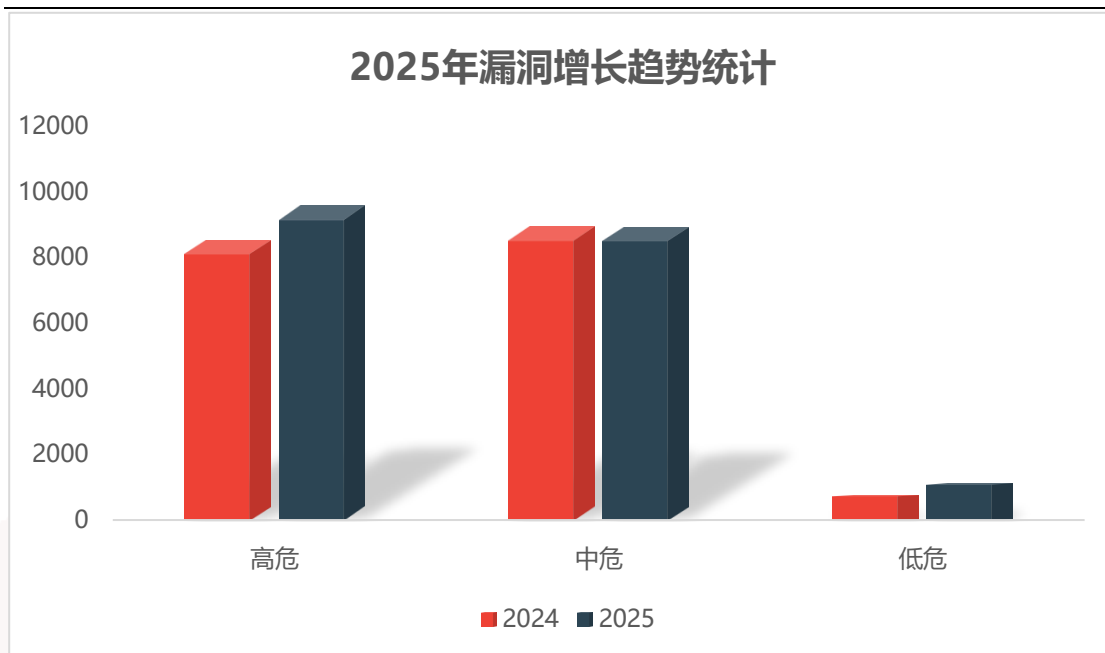


图 8 2025 年漏洞增长趋势统计（数据来自 CNVD 2026.1）

2025 年的漏洞披露情况印证了安全威胁的持续升温与结构演变。结合行业趋势来看，漏洞数量与高危漏洞占比的双升，与攻击者借助 AI 降低攻击门槛、聚焦 API 与聊天机器人等新攻击面的态势高度契合，而低危漏洞的大幅增长则反映出智能设备长生命周期、供应链复杂等带来的长尾风险不断暴露。2026 年，企业需以漏洞数据为预警，一方面强化 AI 驱动的漏洞挖掘与实时响应能力，针对高危漏洞建立快速闭环处置机制，另一方面将漏洞防护延伸至 LLM、API、物联网设备及软件供应链全链条，同时推进 API 全生命周期安全管控与智能体交互的安全治理，以此应对漏洞增长与攻击模式演变的双重挑战，避免 AI 落地、物联网普及等带来的风险敞口持续扩大。

第四章 在野利用漏洞概况

通过对全网漏洞数据的分析统计,我们筛选了 2025 的前 100 个在野利用漏洞进行统计分析。此次统计分析主要从漏洞所影响厂商、影响平台产品、漏洞类型以及 EXP 公开情况等 4 个方面展开。统计显示，漏洞影响厂商前三为 Microsoft、Google、Cisco，攻击者重

点聚焦数字生态基础组件厂商以实现大范围攻击。

受影响平台分六类，操作系统与系统软件占比最高，网络与服务器软件、Web 应用与框架次之，零日攻击重心向基础设施倾斜。

漏洞类型中，输入验证不当居首，反序列化不信任数据及多种漏洞占比显著，此类低门槛漏洞在 AI 赋能下加剧安全风险。

EXP 公开占比 57%、未公开占比 43%，未公开占比上升体现攻击者转向定向渗透，AI 技术降低其对公开 EXP 的依赖。

2025 年在野漏洞凸显安全形势严峻及企业管控短板，本章结合统计结果拆解隐患与防控要点，助力企业构建主动防御体系。

4.1 漏洞影响厂商分布情况

根据 2025 年在野利用漏洞前 100 例所影响的厂商情况进行统计，前三名分别是 Microsoft、Google、Cisco。其中 Microsoft 的产品占比达到 13%，Google 占比为 7%，Cisco 占比为 5%。

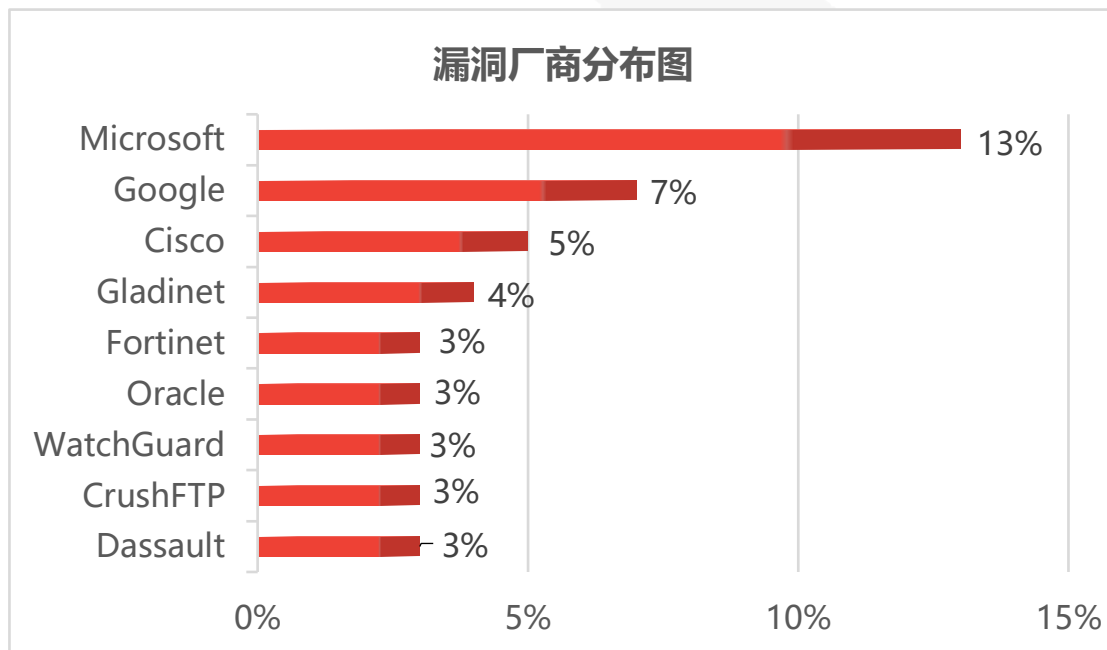


图9 漏洞厂商分布图(数据来自于 cybersecurity-help)

根据 2025 年在野利用漏洞的统计, Microsoft、Google、Cisco 三大科技巨头位列受影响厂商前三名。表明攻击者的策略高度聚焦于数字化生态的“基础组件”。这些厂商的操作系统、浏览器、云服务及网络设备构成了全球数字空间的基石, 拥有最广泛的部署量。因此, 针对它们的漏洞利用能为攻击者带来极高的投入产出比, 实现最大范围的攻击覆盖。在漏洞治理中, 必须对这类无处不在的基础软件和基础设施给予最高级别的优先级和持续的监控。

4.2 漏洞影响平台产品分类

根据 2025 年在野利用漏洞前 100 例所影响的平台产品情况进行统计, 受影响的平台产品大致可分为六类: 分别是操作系统与系统软件、网络与服务器软件、Web 应用与框架、数据库与开发工具、硬件与嵌入式系统以及其他软件。其中占比为操作系统与系统软件 31%、网络与服务器软件 25%、Web 应用与框架 18%、其他软件 12%、硬件与嵌入式系统 10%、数据库与开发工具 4%。

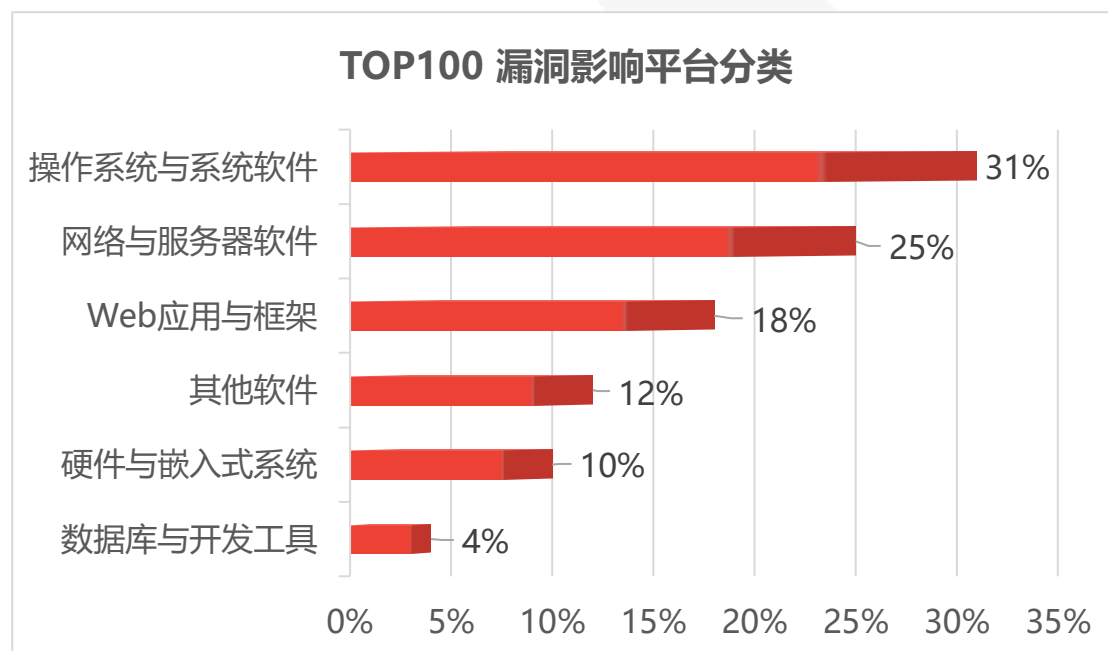


图 10 TOP100 漏洞平台分类(数据来自于 cybersecurity-help)

2025 年数据显示，零日攻击的战略重心发生了变化。从技术进展角度看，这一现象源于三个关键因素。首先，基础设施类漏洞因其前置性特征，能够成为对整个企业网络的立足点，相比单一应用程序的价值倍增。其次，Web 应用框架虽然占比 18%，但其中包含了 React 19 Server Components (CVE-2025-55182, CVSS 10.0) 这类供应链级别的威胁，单一漏洞波及数百万开发者。第三，硬件与嵌入式系统占比虽仅 10%，但包括一些边界设备，其战略位置使其成为 APT 级攻击者的优先目标。

实践启示在于，组织的防御体系应从“应用级补丁”向“网络边界与操作系统基础性防护”倾斜，尤其需要加强对基础设施设备的零日监控与隔离策略，因为这些设备一旦被攻陷，已有的应用层防控将形同虚设。

4.3 漏洞类型统计概况

根据 2025 在野利用漏洞前 100 例漏洞类型情况进行统计，其中输入验证不当 (CWE-20) 占比最多，以 10% 位居首位，反序列化不信任数据 (CWE-502) 占比 7%，路径遍历 (CWE-22) 占比 5%，越界写入 (CWE-787) 占比 5%，缺少授权 (CWE-862) 占比 5%，代码注入 (CWE-94) 占比 4%，操作系统命令注入 (CWE-78) 占比 4%，其余为占比较少的其他类型。

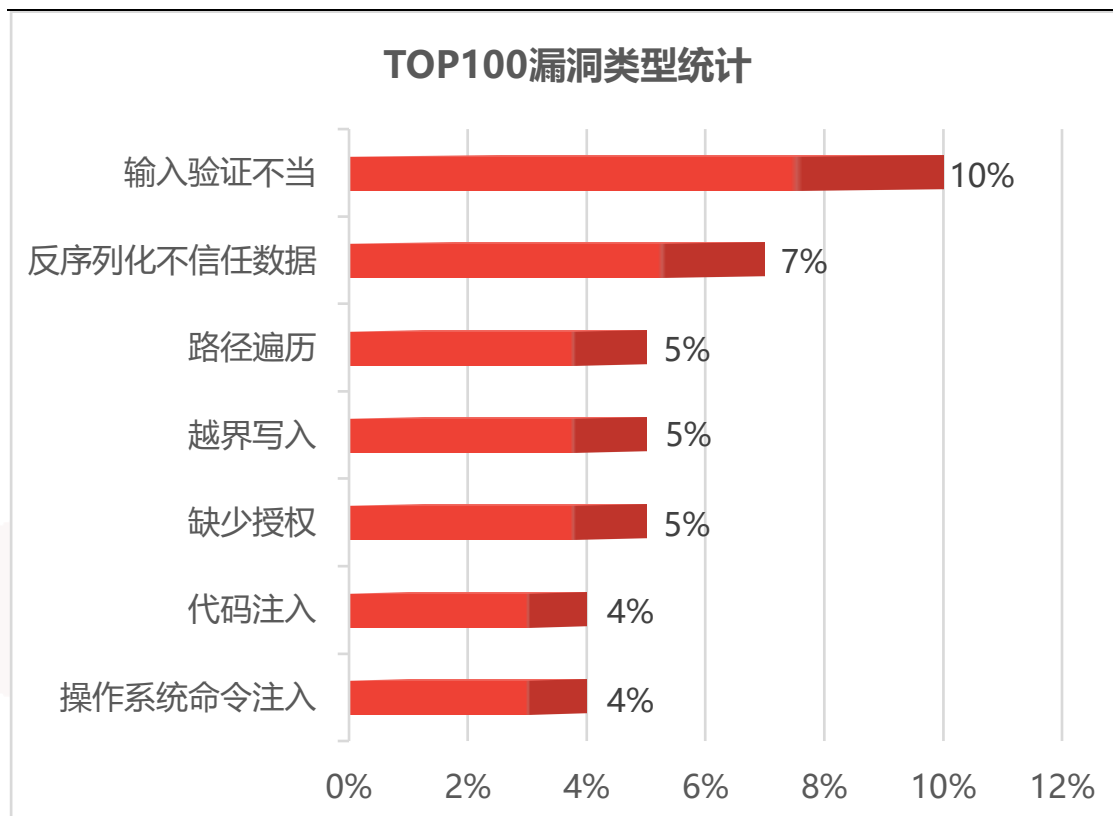


图 11 TOP100 漏洞类型统计概况(数据来自于 cybersecurity-help)

根据 2025 年在野利用漏洞的统计分析，当前网络攻击的核心切入点仍聚焦于基础安全缺陷，输入验证不当（CWE-20）以 10% 占比位居首位，反序列化不信任数据、路径遍历等传统高危漏洞类型也占据显著比例，这与年度“披露即利用”的攻击新常态形成呼应。这类漏洞的共性在于利用门槛低、攻击链成熟，恰好契合了 AI 赋能下攻击者的技术需求：生成式 AI 可快速针对输入验证不当等缺陷编写利用脚本，反序列化漏洞无需认证的特性更被批量转化为自动化攻击工具，成为 APT 组织与勒索团伙实现初始入侵的标配。数据背后也显示出企业基础安全管控的短板，若不能将补丁管理、权限控制等基础措施升级为常态化、自动化的运营能力，即便面对看似简单的传统漏洞，也难以抵御 AI 加持下的规模化、高速化攻击。

根据 MITRE 今年通过对公开可用的国家漏洞数据库中 39080 项数据调研分析得出的 CWE TOP 25 排名如下。

排名	ID	名称	KEV 数量	与 2024 年排名 相比
1	CWE-79	跨站脚本	7	0
2	CWE-89	SQL 注入	4	1
3	CWE-352	跨站请求伪造 (CSRF)	0	1
4	CWE-862	缺少授权	0	5
5	CWE-787	越界写入	12	-3
6	CWE-22	路径遍历	10	-1
7	CWE-416	释放后使用	14	1
8	CWE-125	越界读取	3	-2
9	CWE-78	OS 命令注入	20	-2
10	CWE-94	代码注入	7	1
11	CWE-120	经典缓冲区溢出	0	无
12	CWE-434	危险文件上传	4	-2
13	CWE-476	NULL 指针解引用	0	8
14	CWE-121	栈溢出	4	无
15	CWE-502	反序列化	11	1
16	CWE-122	堆溢出	6	无
17	CWE-863	授权错误	4	1
18	CWE-20	输入验证不当	2	-6
19	CWE-284	访问控制不当	1	无
20	CWE-200	信息泄露	1	-3

21	CWE-306	缺少关键功能的身份验证	11	4
22	CWE-918	服务器端请求伪造 (SSRF)	0	-3
23	CWE-77	命令注入	2	-10
24	CWE-639	越权访问	0	6
25	CWE-770	资源耗尽	0	1

表 1 CVE TOP 25 排名(数据来自于 MITRE)

与过去几年一样，CWE 团队在分析今年的变化时指出，前 25 名的漏洞越来越多地转向内存安全、权限管理和访问控制，在今年的漏洞排行榜上，有几种漏洞类型的排名与去年有所变化，其中有的完全消失或是首次进入前 25 名。

排名大幅提升的漏洞有：

- (1) CWE-862：缺少授权，从第 9 位上升至第 4 位；
- (2) CWE-476：NULL 指针解引用，从第 21 位上升至第 13 位；
- (3) CWE-306：缺少关键功能的身份验证，从第 25 位上升至第 21 位；

排名大幅下降的漏洞有：

- (1) CWE-787：越界写入，从第 2 位下降至第 5 位；
- (2) CWE-20：输入验证不当，从第 12 位降至第 18 位；
- (3) CWE-200：信息泄露，从第 17 位降至第 20 位；
- (4) CWE-918：服务器端请求伪造（SSRF），从第 19 位降至第 22 位；
- (5) CWE-77：命令注入，从第 13 位降至第 23 位；

Top 25 中的新入围的漏洞有：

- (1) CWE-120：经典缓冲区溢出，排名第 11 位，此前排名未上榜；

- (2) CWE-121: 栈溢出, 排名第 14 位, 之前排名未上榜;
- (3) CWE-122: 堆溢出, 排名第 16 位, 之前排名未上榜;
- (4) CWE-284: 访问控制不当, 排名第 19 位, 之前排名未上榜;
- (5) CWE-639: 越权访问, 从第 30 位上升至第 24 位;
- (6) CWE-770: 资源耗尽, 从第 26 位上升至第 25 位;

从 Top 25 中落选的漏洞有:

- (1) CWE-269: 权限管理不当, 从第 15 位下降至第 29 位, 下降 14 位;
- (2) CWE-190: 整数溢出, 从第 23 位下降至第 30 位, 下降 7 位;
- (3) CWE-287: 身份验证不当, 从第 14 位下降至第 31 位, 下降 17 位;
- (4) CWE-400: 不受控制的资源消耗, 从第 24 位降至第 32 位, 下降 8 位;
- (5) CWE-798: 硬编码凭证, 从第 22 位降至第 35 位, 下降 13 位;
- (6) CWE-119: 缓冲区溢出, 从第 20 位降至第 39 位, 下降 19 位;

4.4 EXP 公开情况统计

根据 2025 在野利用漏洞前 100 例 EXP 公开情况进行统计, 其中公开 EXP 稍多, 占比 57%, 未公开 EXP 的为 43%。

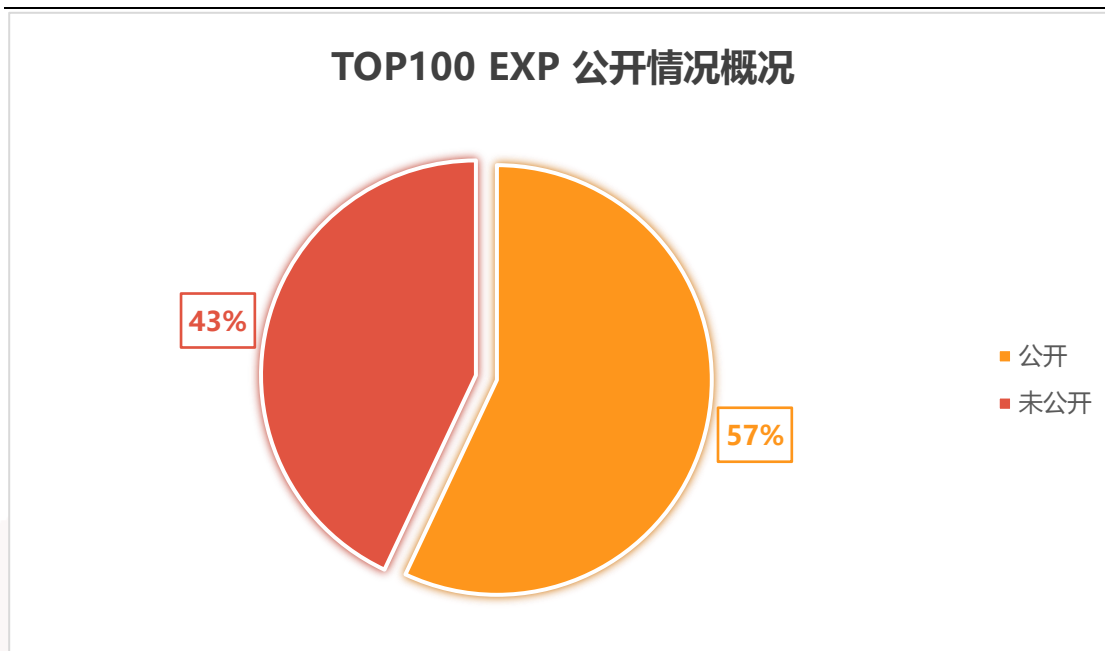


图 12 TOP100 EXP 公开情况概况(数据来自于 cybersecurity-help)

从 2025 年在野利用漏洞前 100 例 EXP 公开情况的统计数据来看，公开 EXP 占比从上年的 72% 降至 57%，未公开占比则从 28% 升至 43%。结合当前漏洞利用“零日化”“武器化”的趋势，未公开 EXP 占比的提升，本质是攻击者（尤其是 APT 组织、勒索团伙）对漏洞利用窗口期的刻意把控，相较于以往公开 EXP 以扩大攻击覆盖面的模式，如今攻击者更倾向于将 EXP 作为定向渗透的专属资源，用于针对关键基础设施、政企核心系统等高价值目标的隐蔽攻击，以此规避防御方通过公开情报快速构建防护规则的可能。

与此同时，AI 驱动的漏洞挖掘与攻击链自动化技术，让攻击者无需依赖公开社区共享，即可快速生成定制化 EXP 并实施精准打击，进一步助推了未公开 EXP 的规模化应用。这种趋势对企业防御体系提出更高要求：单纯依赖公开漏洞情报的被动防护模式已难以为继，必须转向资产全生命周期监测、异常行为实时研判、自动化应急响应的主动防御架构，才能应对未知威胁带来的冲击。

第五章 最值得关注 TOP10 高危漏洞情况

2025 年，天融信阿尔法实验室持续监测全球漏洞态势，累计捕获并分析漏洞情报数万条。通过对海量信息的实时追踪与快速研判，实验室成功预警并协助处置了多起突发性的
高危漏洞安全事件。基于漏洞的影响范围、潜在危害程度及实际威胁等级，系统性地梳理
并发布了年度最具代表性的十大高危漏洞，具体分析情况如下：

5.1 年度 TOP10 高危漏洞

本节内容筛选自天融信 2025 年预警的漏洞信息，并根据漏洞的利用难易程度、漏洞
利用成功后造成的损失、漏洞影响的范围进行排名，根据排名节选出排名前十的漏洞。

危害程 度排名	漏洞编号	标题	概述
NO.1	CVE-2025-55182	React 服务端组件远 程 代 码 执 行 漏 洞 (React2Shell)	React 服务端组件在处理 Flight 协议数据时存在反序列化漏洞， 攻击者可通过特制请求实现远程 代码执行。
NO.2	CVE-2025-33073	Windows SMB 权限 提升漏洞	Windows SMB 协议中存在权限提 升漏洞，攻击者通过恶意 DNS 记 录绕过 NTLM 反射保护，可在未 启用签名的系统上获取 SYSTEM 权限。
NO.3	CVE-2025-59287	Microsoft WSUS 未 授权反序列化漏洞	Microsoft WSUS 因不安全的反序 列化存在漏洞，未授权攻击者可

			通过发送恶意加密 Cookie 以 SYSTEM 权限远程执行代码。
NO.4	CVE-2025-49844 (RediShell)	Redis Lua 脚本解析 释放后重引用 UAF 远程代码执行漏洞	Redis 处理 Lua 脚本时存在 UAF 漏洞，攻击者通过恶意脚本可在服务器上远程执行任意代码。
NO.5	CVE-2025-32463	Sudo 外部资源引用 不当本地权限提升漏洞	Sudo 的 -R 选项在解析路径时存在缺陷，本地用户可通过创建伪造配置文件加载恶意库，从而获取 root 权限。
NO.6	CVE-2025-49704 CVE-2025-49706 CVE-2025-53770 CVE-2025-53771	Microsoft SharePoint ToolShell 漏洞利用 链	SharePoint 中存在多个漏洞（包括反序列化与认证绕过），攻击者可组合利用这些漏洞上传 WebShell 并执行任意代码。
NO.7	CVE-2025-1974	Kubernetes Ingress NGINX 准入控制器 远程代码执行漏洞	准入控制器在校验 Ingress 配置时，未对恶意 NGINX 指令进行有效过滤。攻击者向 Webhook 接口发送伪造请求即可触发代码执行，从而接管控制器 Pod 并利用高权限凭据渗透整个 Kubernetes 集群。
NO.8	CVE-2025-3248	Langflow 代码验证 接口未授权代码注入	Langflow 的代码验证接口缺乏身份验证，未授权攻击者可发送特

		漏洞	制请求执行任意代码。
NO.9	CVE-2025-55177 CVE-2025-43300	苹果设备 WhatsApp 零点击漏洞利用链	攻击者可通过 WhatsApp 同步消息诱导设备加载恶意图像，利用 Apple Image I/O 漏洞实现零点击远程代码执行。
NO.10	CVE-2025-8088	WinRAR 路径穿越漏洞	WinRAR 早期版本存在路径穿越漏洞，攻击者可通过构造的压缩包将恶意文件释放至系统启动目录，实现代码执行。

表 2 TOP10 危害排名

5.2 年度漏洞预警 TOP10 漏洞回顾

一、 React 服务端组件远程代码执行漏洞（React2Shell）
漏洞类型：远程代码执行
漏洞编号：CVE-2025-55182
CVSS:3.1/AV:N/AC:L/PR:N/UI:N/S:C/C:H/I:H/A:H
漏洞评分：10.0
预警日期：2025-12-04
漏洞描述：该漏洞被安全社区命名为“React2Shell”，是一个位于`react-server`包中的严重反序列化漏洞。该漏洞的根源在于 React 处理 Flight 协议数据流的方式。 当使用 Next.js App Router 的应用接收到客户端请求时（通常是 POST 请求），服务器端需要解析请求体中的序列化数据以执行 Server Actions 或重构组件状态。由于

`react-server`在解析过程中未能对输入数据进行严格的类型验证和结构检查，攻击者可以构造一个包含恶意对象的特制 Payload。这些恶意对象利用了 JavaScript 的原型链机制 (Prototype Pollution)，通过注入特定的属性（如`\_\_proto\_\_`或特定的内部状态标记），攻击者能够覆盖服务器端的全局对象方法。

二、Windows SMB 权限提升漏洞

漏洞类型：权限提升漏洞

漏洞编号：CVE-2025-33073

CVSS:3.1/AV:N/AC:L/PR:L/UI:N/S:U/C:H/I:H/A:H

漏洞评分：8.8

预警日期：2025-06-16

漏洞描述：SMB (Server Message Block) 是一种网络文件共享协议，主要用于在局域网中实现文件、打印机和串行端口的共享。SMB 协议是一种客户端-服务器架构的请求/响应协议，客户端应用程序可通过该协议在各种网络环境下读写服务器上的文件，并对服务器程序提出服务请求。

该漏洞通过特制的 DNS 记录（如 "srv11UWhRCAAAAAAAAAAAAAAAAAAAAAAAAAAwbEAYBAAAA"）强制目标机器进行本地 NTLM 认证，从而绕过 SMB 到 SMB 的 NTLM 反射保护机制。攻击者可以利用恶意脚本迫使受害机器通过 SMB 协议连接回攻击者系统并进行身份验证，在客户端和服务端认证过程中利用逻辑判断错误，最终在未强制执行 SMB 签名的机器上以 SYSTEM 权限执行任意命令。

三、Microsoft WSUS 未授权反序列化漏洞

漏洞类型：远程代码执行

漏洞编号：CVE-2025-59287

CVSS:3.1/AV:N/AC:L/PR:N/UI:N/S:U/C:H/I:H/A:H

漏洞评分：9.8

预警日期：2025-10-27

漏洞描述：WSUS（Windows Server Update Services，Windows 服务器更新服务）是微软提供的一项服务器角色，用于在企业内部网络中集中管理、分发和安装 Windows 操作系统及其他微软产品的更新。

该漏洞是由于 EncryptionHelper.DecryptData()方法使用 BinaryFormatter 类对 AuthorizationCookie 对象数据进行了不安全的反序列化，未经身份验证的远程攻击者可以向 GetCookie()端点发送恶意加密 Cookie 来利用此漏洞，从而允许攻击者以 SYSTEM 权限远程执行代码。

四、Redis Lua 脚本解析释放后重引用 UAF 远程代码执行漏洞

漏洞类型：远程代码执行

漏洞编号：CVE-2025-49844

CVSS:3.1/AV:N/AC:L/PR:L/UI:N/S:C/C:H/I:H/A:H

漏洞评分：9.9

预警日期：2025-11-03

漏洞描述：Redis 是一个开源的内存键值数据库，支持数据持久化到磁盘。其广泛应用于缓存、消息队列、会话存储等场景，在互联网公司和企业级应用中占据重要地位。

Redis 处理 Lua 脚本的代码中存在一个释放后重引用 UAF 漏洞，经过身份验证的远程攻击者可通过精心构造的 Lua 脚本触发此漏洞，这会导致某些内存区域被提前释放，但相关指针仍然被后续代码引用，从而触发 UAF。成功利用此漏洞的攻击者可以在 Redis 服务器上远程执行任意代码。此漏洞由 Pwn2Own Berlin 2025 的参赛者发现，其已在 Redis 源代码中存在约 13 年。

五、Sudo 外部资源引用不当本地权限提升漏洞

漏洞类型：本地权限提升

漏洞编号：CVE-2025-32463

CVSS:3.1/AV:L/AC:L/PR:N/UI:N/S:C/C:H/I:H/A:H

漏洞评分：9.3

预警日期：2025-06-30

漏洞描述：Sudo 是 Linux/Unix 中的一个工具，其允许系统管理员委派权限，使特定用户或用户组能够以 root 用户或其他用户身份运行部分或全部命令，同时提供命令及其参数的审计跟踪。

Sudo 的-R (--chroot) 选项旨在允许用户在 sudoers 文件允许的情况下，使用用户指定的根目录运行命令。Sudo 1.9.14 版本进行了一项更改，在 sudoers 文件仍在解析时，chroot()函数会使用用户指定的根目录来解析路径。本地低权限用户可通过在用户指定的根目录下创建/etc/nsswitch.conf 文件来欺骗 Sudo 加载任意共享库，从而直接获取超级用户 (root) 权限。

六、Microsoft SharePoint ToolShell 漏洞利用链

漏洞类型：反序列化不安全数据

漏洞编号：CVE-2025-49704、CVE-2025-49706、CVE-2025-53770、CVE-2025-53771

CVSS:3.1/AV:N/AC:L/PR:N/UI:N/S:U/C:H/I:H/A:H

漏洞评分：9.8

预警日期：2025-07-20

漏洞描述：Microsoft SharePoint 是微软公司开发的企业级协作和文档管理平台，广泛应用于企业内部的文档共享、团队协作、工作流管理等场景。作为微软 Office 365 和 SharePoint Server 的核心组件，SharePoint 在全球企业级市场占据重要地位，为数百万用户提供文档库、列表、网站集合等功能。

CVE-2025-49704 是由于 Microsoft SharePoint 对 XML 内容验证不当导致的反序列化漏洞，经过身份认证的远程攻击者可通过此漏洞上传 WebShell，进而窃取加密密钥，最终在受感染的服务器上执行任意代码。

CVE-2025-49706 是 Microsoft SharePoint 中的身份认证绕过漏洞，未经身份认证的远程攻击者可通过伪造 Referer 头来向/ToolPane.aspx 页面发送精心构造的 POST 请求，使服务器误以为该请求来自应用程序内部，从而绕过正常的身份认证检查，并允许攻击者访问特权功能。

CVE-2025-53770 是 CVE-2025-49704 的补丁绕过，CVE-2025-53771 是 CVE-2025-49706 的补丁绕过。这些漏洞由 Pwn2Own Berlin 2025 的参赛者发现，组合利用这些漏洞形成 ToolShell 漏洞利用链。

七、Kubernetes Ingress NGINX 准入控制器远程代码执行漏洞

漏洞类型：远程代码执行

漏洞编号: CVE-2025-1974

CVSS:3.1/AV:N/AC:L/PR:N/UI:N/S:U/C:H/I:H/A:H

漏洞评分: 9.8

预警日期: 2025-03-25

漏洞描述: Ingress 是 Kubernetes 的传统功能, 用于将工作负载 Pod 公开给外界, 以便它们发挥作用。Ingress-nginx 是 Kubernetes 项目提供的纯软件入口控制器。由于其多功能性和易用性, 在超过 40% 的 Kubernetes 集群中都部署了 Ingress-nginx。

准入控制器在校验 Ingress 配置时, 未对恶意 NGINX 指令进行有效过滤。攻击者向 Webhook 接口发送伪造请求即可触发代码执行, 从而接管控制器 Pod 并利用高权限凭据渗透整个 Kubernetes 集群。

八、Langflow 代码验证接口未授权代码注入漏洞

漏洞类型: 远程代码执行

漏洞编号: CVE-2025-3248

CVSS:3.1/AV:N/AC:L/PR:N/UI:N/S:U/C:H/I:H/A:H

漏洞评分: 9.8

预警日期: 2025-04-08

漏洞描述: Langflow 是一款用于构建多代理和 RAG (检索增强生成) 应用程序的开源可视化低代码框架。在 1.3.0 之前的版本中, 系统的 validate_code 验证组件存在严重的安全缺陷。

该漏洞的核心在于 Langflow 的 /api/v1/validate/code 端点 (或在某些版本中为 /api/v1/builder/execute_code) 未实施任何身份验证机制。攻击者可以向该接口发送

包含恶意 Python 代码的 HTTP POST 请求。由于系统在后端直接调用不安全的 `exec()` 函数或 `compile()` 函数来解析和执行用户提交的 Python 代码（包括滥用 Python 装饰器和默认参数等高级语法），攻击者无需登录即可在服务器操作系统上下文中执行任意系统命令。

九、苹果设备 WhatsApp 零点击漏洞利用链

漏洞类型：越界写入与授权绕过

漏洞编号：CVE-2025-55177、CVE-2025-43300

CVSS:3.1/AV:N/AC:L/PR:N/UI:N/S:C/C:H/I:H/A:H

漏洞评分：9.8

预警日期：2025-08-29

漏洞描述：WhatsApp 是全球广泛使用的端到端加密即时通讯应用。该漏洞涉及两个关键 CVE 的零点击利用链：WhatsApp 的 CVE-2025-55177 和苹果 iOS/macOS 的 CVE-2025-43300。CVE-2025-55177 源于 WhatsApp 链接设备同步消息授权不完整。攻击者无需用户交互，通过精心构造的同步消息，诱导目标设备从任意 URL 处理恶意内容。CVE-2025-43300 是苹果 Image I/O 框架中的越界写漏洞，发生在 DNG 图像格式的解压缩代码中，攻击者可制造恶意 DNG 图像触发内存损坏。

十、WinRAR 路径穿越漏洞

漏洞类型：路径穿越

漏洞编号：CVE-2025-8088

CVSS:3.1/AV:N/AC:L/PR:N/UI:R/S:U/C:H/I:H/A:H

漏洞评分：8.8

预警日期：2025-08-12

漏洞描述：WinRAR 是一款流行的文件压缩和解压缩软件，由 RARLAB 开发和维护。它支持多种压缩格式，包括 RAR、ZIP、ISO 等。WinRAR 在全球范围内拥有大量用户，是 Windows 平台上最常用的压缩软件之一。

WinRAR 7.13 之前的版本中，由于其对压缩包内文件路径的校验机制存在缺陷，当用户解压文件时，程序可能错误地使用压缩包内预设的路径，而非用户指定的目标路径。通过植入特殊构造的相对路径（如..\），攻击者可诱使程序将恶意文件释放至系统启动目录或其他敏感位置，最终实现远程代码执行。

第六章 2026 年漏洞技术走向

预计在 2026 年，全年漏洞总量仍将保持现有水平甚至会出现一定幅度的增长。保持现有水平的主要依据是，现有开发及安全技术无重大变革，智能编程作为一种辅助开发手段，仍在迭代中。出现小幅增长的主要依据是，AI 模型及应用持续迭代，相关的基础设施漏洞会是新的爆发点，也是新的焦点。一定会引入和发现相当的安全漏洞。技术继续发展，我们预测 2026 年漏洞技术走向存在以下趋势：

- （1）基于 AI 大模型的智能编程模式继续迭代，如何提升产出代码的安全性，是值得关注的重要内容；
- （2）供应链相关的漏洞，仍会是攻击者亲睐的目标，也是防守方应该重视的地方；
- （3）安全与 AI 的结合将更加紧密，攻击方利用 AI 进行安全测试、漏洞挖掘、漏洞利用的研究将变为日常，防守方深度结合 AI 进行漏洞修复、告警降噪、流量识别等能力或将

更加精准有效。随着 AI 自身能力与流程度（日活量）的进一步提升，攻防速度一定会进一步加快。

大模型提供的网络服务，涉及大量基础设施建设，夹杂模型自身安全风险，其将成为新的攻击热点。

6.1 “智能编程”安全漏洞不可避免

在软件开发领域，AI 大模型正推动智能编程进入一个全新的时代。回顾 2024 年，许多开发者已开始借助大模型的代码生成能力，将其生成的函数级代码手动整合到项目中，实现初步的辅助编程。这一阶段的 AI 编程，仍停留在小规模、片段化的代码生成与应用。

进入 2025 年，随着“深度思考”机制的显著进步，AI 大模型的推理能力大幅增强，推动更智能的 Vibe Coding 模式走向现实。在这一模式下，开发者无需逐行编写代码或强记语法，而是直接通过自然语言（如中文、英文）向 AI 发出指令，依托 VS Code、Cursor 等集成开发环境，由大模型自动完成编码、测试，并输出可直接运行的软件产品。借助更强大的上下文理解与复杂逻辑处理能力，Vibe Coding 已能应对真实软件工程需求，而不仅仅是编写演示原型。这种模式的演进，显著提升了开发效率，以往需数小时完成的任务，如今可在几分钟内解决。

Stack Overflow 于 2025 年 11 月发布的《2025 年度开发者调查报告》显示，共有来自 177 个国家的超 49,000 名开发者参与，涵盖 314 项技术相关议题。调查指出，2025 年有 84% 的受访者正在使用人工智能工具。

GitHub Copilot 是由 GitHub 与 OpenAI 共同开发的人工智能编程助手，能够在编程时提供实时代码补全与建议，帮助开发者提升编码效率。它通过分析海量代码库学习编程模式，并已集成到 VS Code、Visual Studio、JetBrains IDEs 等主流编辑器中，能够

依据上下文智能生成函数、代码块乃至测试用例。

微软作为 GitHub 的母公司，会在财报中披露其核心产品数据。2025 年中期数据显示，GitHub Copilot 用户数（含试用与付费）已突破 2000 万，且超过 90% 的财富 100 强企业已使用该产品。作为背景参考，据 Evans Data / Statista 2024 年统计，全球专业软件开发者数量约为 2870 万人。

在此背景下，当智能编程模式的日活率进一步提升，代码控制力下降叠加大量的 AI 生成代码，软件产品的安全漏洞态势不容乐观。

6.2 供应链的漏洞攻击仍威胁严重

供应链相关的漏洞攻击事件数量之所以长期居高不下，主要是因为软件的生产模式及攻击者的投入产出比两方面，使得情势朝着这个方向不断发展。这不再是单纯的“Bug 修复”问题，而是一场不对称的攻击。

只要软件开发模式依然依赖第三方库及组件，供应链风险就永远存在。现在的软件开发模式，软件不是“写”出来的，而是“组装”出来的，可以说 90% 的代码源自第三方组件。你引入了一个简单的日志库（比如 log4j），它可能又依赖了另外 5 个库，这 5 个库又依赖了 20 个库。结果就是你看不见的“间接依赖”呈指数级膨胀。只要最底层的某个不起眼的小组件出问题，整个金字塔尖的数万个应用都会受到影响。

攻击者也是理性的，供应链攻击是目前投资回报率最高的攻击手段。传统模式下，攻击一个系统，需要研究它的架构、代码，费时费力攻破后只能影响这一个系统。在供应链模式下，攻破一个像 React、Log4j 或 npm 这样的系统，一次投入，就能影响成千上万个系统。

回顾 React2Shell 事件，未来的供应链攻击不一定是植入恶意代码，而是挖掘合法代

码中的深层逻辑缺陷。以利用供应链的广泛影响，扩大影响范围。随着 Next.js、Nuxt、Spring 等全栈框架越来越复杂，框架层面的反序列化、渲染逻辑漏洞将成为重灾区。2026 年，基于供应链的漏洞攻击仍将是软件产业的严重威胁，应警惕在关键时刻针对核心设施或组件的供应链攻击，把防治供应链漏洞视为安全运维的重要目标。

6.3 AI 漏洞挖掘与利用效率将进一步提升

2025 年的情况显示，AI 大模型的应用显著降低了攻击门槛，使得漏洞发现与利用更为便捷，截至 2026 年初，利用 AI 进行漏洞挖掘与利用已经从概念走向了部分成熟的实战应用阶段，但两者成熟度存在显著差异，漏洞挖掘已较为成熟并融入工业界流程，而自动编写漏洞利用仍处于“人机协同”或特定场景（如 CTF、简单漏洞）有效的阶段。

在漏洞挖掘方面已经有成熟的商业和开源应用。AI 不再仅仅是辅助，Google 的 OSS-Fuzz 已集成大型语言模型来自动生成模糊测试用例，覆盖率远超传统方法。2025 年初，Google 的 AI Agent “Big Sleep” 成功在 SQLite 等知名开源项目中完全自主发现了 0-day 漏洞，这是 AI 漏洞挖掘历史上的里程碑。在漏洞利用方面，目前 AI 尚未达到“一键通杀”的成熟度。AI 擅长生成简单的 POC，但在面对需要多步推理、绕过复杂防护或构建复杂利用链时，仍需人工介入。

我们可以从两方面看待 AI 对于漏洞挖掘和利用的帮助。第一方面是辅助安全研究，使用 AI 模型的问答能力，辅助进行代码审计、利用编写，已经可以有效提高人工效率。未来随着模型的能力进化，准确性、效率会更进一步提升。第二方面，使 AI 完全自动进行漏洞挖掘，部分安全产品已具有相关能力，未来随着大模型不断进化，相关能力势必稳步提升。关于 AI 自动漏洞利用，鉴于漏洞利用的复杂性、知识匮乏等问题，则将是未来需要重点突破的地方。

6.4 大模型自身的漏洞将成为新的攻击热点

2025 年的情况还显示出，新兴的大模型技术本身也成为了新的安全测试目标。大模型及其支撑体系正在构成一个新兴的攻击面，上层应用面临提示词注入、生成式 AI 输出验证不当等新型风险，属于 AI 大模型特有类型的漏洞；底层基础设施则依然存在未授权访问、命令执行等传统漏洞。从统计数据上看，AI 及其组件相关漏洞占比已经有一定量的上升。由于这一领域尚未被充分研究，安全体系尚处于建设阶段，值得安全研究人员持续关注与投入，未来也一定是安全研究人员的一个主要目标。

按照黄仁勋提出的五层分层理论将 AI 解构，构建 AI 大模型，可以分成：能源、芯片、基础设施、模型、应用共计五个层面，可以说各个层面都存在其特殊的安全风险，也都是值得关注的方向。我们抛开能源与芯片，重点讨论上面三层存在的安全漏洞风险，以下是一些可供研究人员参考的方向。

危害程度排名	层级描述	层级漏洞特点	典型漏洞案例
应用	第五层，最上层的应用层，调用下层模型，为用户提供服务的接口。	大模型特有漏洞与传统漏洞均涉及。大模型特有漏洞如：提示词注入、不安全的插件与工具调用、输出处理不当、敏感信息过度披露、拒绝服务	Langflow RCE (CVE-2025-3248)：Langflow 是一个流行的可视化的 AI Agent 构建工具。2025 年，它被发现存在一个低级的未授权远程代码执行漏洞。攻击者无需登录，只需向 /api/v1/validate/code 接口发送一段包含恶意 Python 代码的 HTTP 请求，Langflow 就会乖乖执行。

		等。传统的安全漏洞，均被覆盖。	
模型	第四层，从本质上讲，AI 模型是一个巨大的数学函数，或者是对人类知识的概率压缩。它的核心功能是预测，现代大模型几乎都基于 Transformer 架构。	模型特有安全漏洞，如数据投毒与后门、模型文件格式漏洞、隐私泄露漏洞、模型窃取。	Nightshade (艺术家反击/数据投毒): 为了抗议 AI 公司未经许可抓取画作训练，芝加哥大学团队开发了 Nightshade 工具。它通过对图片像素进行肉眼不可见的微调 (添加对抗性扰动)，让 AI 模型在训练时产生认知错乱。如果一个模型使用了大量经过 Nightshade 处理的“狗”的照片进行训练，模型最终可能会把“狗”识别成“猫”或者“烤面包机”。这本质上是一种数据投毒。攻击者污染了训练数据，导致模型中毒，出现错误认知。
基础设施	第三层，为大模型提供运行环境的基础设施，如服务器环境、数据库、HTTP 服务器等，与传	均为传统漏洞。如：命令执行、XSS、认证问题、授权问题、代码执行问题、文件处理问题等。	Ollama RCE (CVE-2024-37032): 一个经典的路径遍历漏洞。Ollama 允许从互联网拉取模型，攻击者构建了一个恶意的 HTTP 请求，欺骗 Ollama 的 API，使其不仅下载模型，还能覆盖系统上的任意文件。虽然官方已修复，但在 2025 年初的扫

	统基础设施无异。		描中，全球仍有大量未更新的旧版本 Ollama 暴露在公网，成为僵尸网络的首选目标。
芯片	暂不关注这部分漏洞风险，但这部分仍然存在风险，且不可忽视。		
能源			

表 3 AI 五层分层理论

第七章 总结

2025 年网络空间安全漏洞态势分析研究报告显示，AI 驱动下的智能编程、供应链攻击与大模型自身安全已构成核心挑战，漏洞总量及高危比例持续上升，攻击焦点加速向关键基础设施与网络基础组件迁移。通过对 2025 年前 100 个在野利用漏洞统计看出，操作系统、网络设备等基础设施类漏洞占比显著增加，公开 EXP 比例下降，攻击正趋于定向化与武器化，基础安全缺陷如输入验证不当、反序列化问题仍是主要突破口。从年度最具代表性的十大高危漏洞来看，影响广泛的供应链漏洞（如 React2Shell）、关键系统权限提升漏洞（如 Windows SMB）及零点击利用链（如 WhatsApp 漏洞组合）等集中暴露了主流开源框架、企业核心业务系统及 AI 应用基础设施在安全设计与权限管控上的深层短板。

总结 2025 年漏洞技术趋势，首先，智能编程的普及虽提升效率却削弱了代码控制力，叠加 AI 生成代码的广泛使用，安全风险持续积累。其次，供应链攻击依托开源生态获得更强潜伏性与针对性，单一底层漏洞可引发行业级灾难。同时，AI 大模型不仅成为新型攻击面，其提示词注入、训练数据泄露等特有风险开始显现。此外，AI 极大提升了攻防双方效率，在漏洞挖掘、利用脚本编写等方面已深度融入安全流程。最后，网络攻击呈现战略化转向，攻击者更多利用漏洞链针对电力、通信、交易所等关键基础设施实施物理破坏与持

久控制。

展望 2026 年漏洞技术走向，智能编程模式将继续演进，确保 AI 生成代码的安全性将成为重要议题；供应链漏洞因其高回报率，仍将是攻击者重点目标，防守方需强化成分管理与攻击面收敛。AI 与安全的结合将更为紧密，攻击方利用 AI 进行自动化漏洞挖掘与利用的能力会进一步增强，防守方则需借助 AI 实现更精准的威胁检测与响应。与此同时，大模型、大模型应用及其基础设施层将继续暴露出大量传统与新型安全漏洞，成为安全研究的新热点。总体而言，随着 AI 技术渗透与基础设施数字化深化，漏洞威胁将更侧重质量与时效，呈现“披露即利用”的常态，推动防御体系向实时监测、自动化响应与供应链全周期治理升级。

7.1 安全防护建议

在归纳 2025 年网络空间安全态势的过程中，我们也针对态势发展的新形势，提出了一些防护建议。

(1) 构建 AI 驱动的漏洞闭环治理体系：针对 2025 年漏洞“披露即利用”的闪电化攻击特征，企业需摒弃传统“定期扫描 + 人工研判”的滞后模式，构建以 AI 为核心的漏洞全生命周期闭环治理体系。借助生成式 AI 与威胁情报引擎的联动，实现对漏洞情报的实时抓取、自动化分级，基于资产关联性、CVSS 评分、在野利用情况等维度，快速区分“需立即响应的高危漏洞”与“低风险可控漏洞”，解决海量漏洞情报筛选难的痛点。同时，部署 AI 自主修复智能体，针对通用组件漏洞自动生成补丁适配方案，对核心业务系统漏洞则输出精准的临时缓解策略。此外，建立漏洞处置效果的 AI 复盘机制，通过模拟攻击验证修复有效性，形成“情报感知 - 分级处置 - 修复验证 - 策略优化”的闭环，将漏洞响应时间从“日级”压缩至“小时级”，应对攻击窗口持续收窄的挑战。

(2) 打造智能体时代的全域边界防护网：面对智能体渗透业务流程、API 成为攻击主战场的新形势，企业需突破传统边界防护思维，打造覆盖 “网络边界 - API 接口 - AI 组件” 的全域防护网。在网络层，升级下一代防火墙与零信任访问控制系统，针对智能体的自动化访问行为建立动态信任评估模型，区分 “合法业务智能体” 与 “恶意攻击智能体” 的访问特征，实现基于行为的精准拦截。在 API 安全层面，部署专门的 AI API 防护网关，实时监控 LLM 与业务系统间的数据传输，识别异常调用频率、敏感数据泄露等风险，同时建立 API 资产清单与敏感数据映射关系，落实最小权限原则。针对大模型支撑框架等 AI 组件漏洞，需将其纳入核心资产防护范畴，定期开展针对模型训练数据污染、提示词注入等新型攻击的专项检测，同步推动 AI 供应商落实 SBOM（软件物料清单）披露机制，从供应链源头降低风险。

(3) 开展面向智能攻防的全员安全能力重塑：随着黑客智能体的普及，攻击门槛大幅降低，社会工程攻击的目标也从人工客服延伸至 AI 聊天机器人，这要求企业重塑全员安全能力体系，适配智能攻防时代的新需求。一方面，升级安全培训内容，新增智能体攻击防护专项课程，包括识别针对 AI 聊天机器人的诱导式攻击、防范通过 “AI 辅助编码” 引入的漏洞风险、掌握智能体操作的权限边界等内容，通过模拟黑客智能体的攻击场景，让员工直观感受新型威胁的危害。另一方面，建立 “人类员工 + 安全智能体” 的协同防御模式，将普通员工培养为 “安全智能体指挥官”，使其掌握调用安全智能体进行日常风险巡检、异常行为上报的技能，同时明确人机协同的责任边界。此外，强化针对关键岗位的供应链安全意识培训，使其能够识别第三方智能工具的潜在风险，形成 “全员懂智能攻防、全员会协同防御” 的新型安全文化。

(4) 构建动态信任为核心的身份安全防御体系：针对黑灰产 Agent 化攻击下身份威胁凸显的难题，需摒弃 “事后验证身份” 的传统模式，转向 “主动建立信任” 的身份安全

架构，抵御攻击者利用身份层面薄弱点开辟的入侵路径。一方面，融合生物特征、行为模式、环境信息等多维数据，依托垂类 AI 大模型构建动态风险感知能力，精准识别身份冒充、权限滥用等风险，提升复杂场景下风险判断的敏锐度与准确性。另一方面，建立全生命周期数字身份管理机制，将 AI 聊天机器人等智能体纳入身份管控范畴，明确其权限边界与访问规则，防范针对智能体的身份诱导攻击。同时，结合零信任架构理念，实现身份权限的动态适配与最小权限管控，联动行为审计系统，对异常身份操作进行实时告警与拦截，从源头遏制身份攻击风险。

(5) 搭建供应链全链条安全防护体系：面对多点、多层次链式演进的供应链威胁，需聚焦开源套件、AI 模型库等核心环节，构建“源头管控-过程监测-应急响应”的全链条防护体系。在源头层面，推动 AI 供应商及第三方服务商落实 SBOM（软件物料清单）披露机制，建立供应商安全评级与准入制度，对供应链组件开展前置安全检测，从源头降低风险。过程管控中，将 AI 模型支撑框架、开源组件等纳入核心资产防护范畴，定期开展针对模型训练数据污染、供应链组件漏洞的专项检测，强化开源生态安全治理，弥补开源套件脆弱性短板。同时，部署供应链安全监测平台，实现对供应链组件全生命周期的持续监控，建立异常行为基线，快速识别双重漏洞串联等攻击行为。此外，制定供应链安全应急响应预案，明确攻击溯源、隔离处置、替代方案等流程，形成供应链安全防护闭环。

(6) 部署前置式 AI 协同主动防御体系：应对 AI 驱动攻击自动化带来的速度与复杂性挑战，需彻底摒弃传统事后防御思维，构建“预测-拦截-协同-优化”的前置式主动防御体系，践行“用 AI 反 AI”核心思路。部署 AI 驱动的自动化防御系统与预测性威胁防御平台，基于 MITRE ATT&CK 人工智能框架，结合可观测性方法监控 AI 智能体全流程活动，提前预判攻击路径与意图，将防御关口从“攻击后响应”前移至“攻击前预警”。搭建“安全智能体集群+集中式管控平台”架构，实现多智能体协同防御，通过集中式监测完成攻

击行为的快速定位与隔离，将防御响应效率与精准度提升至新层级。同时，建立防御策略动态优化机制，结合攻击态势变化与防御效果复盘，持续迭代 AI 防御模型，适配 AI 攻击手段的快速演进，应对攻击门槛降低、攻击模式升级的挑战。

(7) 筑牢法规合规与产业生态双支撑体系：立足“新质网安”时代要求，以新修订《中华人民共和国网络安全法》为遵循，构建法规合规与产业生态协同发力的安全支撑体系，推动安全能力与新质生产力发展同频。在合规层面，对照新法中关键信息基础设施保护、AI 风险监管等要求，梳理核心业务合规痛点，建立 AI 安全合规评估机制，落实网络安全事件报告管理制度，确保安全建设与法规要求无缝衔接。在产业生态层面，聚焦“国产化、行业化、服务化、智能化”四化融合方向，布局国产化 AI 安全平台，推动行业纵深化防御方案落地，创新服务化安全交付模式。同时，加强 AI 安全人才培养，构建适配智能攻防时代的人才梯队，强化国际合作与信息共享，联合产业链各方建立 AI 安全治理共识。

数字经济与 AI 技术的快速发展给网络安全带来了新的挑战与机遇，网络安全防护需求日趋复杂多元化。天融信坚持“融智跃迁”战略，将人工智能能力贯穿于安全设计、防护部署、模型训练与应用全生命周期，已构建起“框架—产品—服务”的全链路防护体系。公司以总部云服务平台为基石，依托遍布全国的监测能力中心，整合云端专家与线下服务人员的协同优势，通过天问大模型、AI 防火墙、大模型安全网关等创新产品，为各行业客户提供覆盖环境、数据、模型、应用与供应链的“网络+数据+AI”立体化安全防护与运营管理服务。