

2025 守内安信息科技 & ASRC

第三季度邮件安全观察



2025 年第三季度, 电子邮件攻击呈现出“手法更精细化、攻击链更隐蔽、更善于利用合法服务作为跳板”的发展趋势。攻击者广泛运用人工智能工具强化社交工程手段, 不仅使邮件文案的说服力显著提升, 还使其在多语言场景下的识别难度大幅增加; 与此同时, 攻击者擅长借助短链接、云平台、电子签章以及安全厂商自身的威胁链接替换机制等第三方服务, 实现攻击路径的串联、跳转与伪装, 导致传统基于静态过滤和来源信誉的防护机制全面失效。

针对网络安全厂商的社交工程试探行为同样明显增多。攻击者通过“少量多次”的策略, 在厂商公开服务接口植入恶意软件或开展情报搜集活动, 试图构建可长期潜伏的后门通道。

最终, 人工智能技术的普及不仅提高了钓鱼攻击的成功率, 还推动漏洞挖掘、自动化探测以及隐蔽指令滥用等攻击行为的效率提升。这一变化使得电子邮件端到端防护必须完成从单层规则防御到整合式防护体系的升级, 该体系需涵盖行为与情境感知、跨渠道监控以及数据防泄漏等核心能力。

以下为 ASRC 与守内安在这一季的特殊观察:

恶意链接置换功能遭到滥用

为降低收件者误点恶意链接的风险, 许多邮件防护机制在邮件传递途中执行「**置换链接 (URL Rewriting)**」, 将邮件内原始链接改写为防护服务自有的检查跳板, 以便在使用者点击时实时检查链接最终的安全性, 并记录用户与点击时间, 以便于事后事后甄别或封锁。这种实时检测机制在传统钓鱼攻击中能有效提高拦截率及方便事后威胁溯源。

然而在 2025 年第三季度, 出现明显滥用的趋势: 攻击者串接多个置换链接与合法跳转服务 (例如短链接、合法云端或第三方追踪域名), 形成「多级跳转」的攻击链。其操作逻辑与风险如下:

1. 串接防护跳板以躲避实时检测

攻击者先利用合法服务 (或被入侵的服务) 生成短链接或跳转链接, 再将这些链接放入钓鱼邮件中。当收件者点击时, 第一个被检查到的 URL 可能是某网安厂商或其它合法防护的置换域名, 因其来源被视为「可信」, 系统就不会进一步深度解析或标示为可疑, 导致最终恶意链接得以通过。

2. 绕过记录与追踪机制

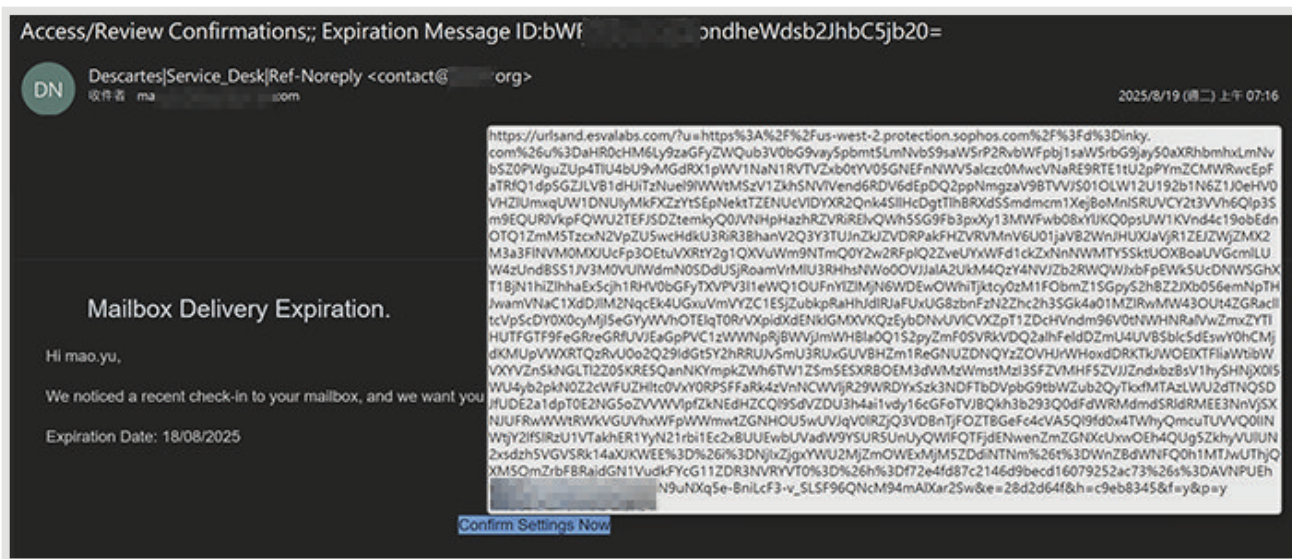
若多段跳转使中间某些 Click-tracking / 置换节点被系统视为正常流量, 系统可能不会完整纪录最终网址或点击人信息, 削弱事后甄别能力与事后责任归属。

3. 利用合法资源作掩护

当跳转链接包含受信任的第三方 (例如广告追踪、电子签章或大品牌云端), 攻击行为将显得更「自然」, 让用户与系统更难辨认其恶意意图。

4. 自动化与规模化

攻击者可以自动化生成大量多段跳转链接, 配合 AI 编写的拟人化话术, 钓鱼效率得到显著提升。



▮ 钓鱼邮件的攻击者尝试串起不同防护的跳转链接，并搭配缩写、跳转的功能，让置换链接防护失效

潜在影响

防护网关的“第一层检查”易被攻击的合法外壳误导，导致“伪阴性”误判率显著上升；事后追溯时鉴别信息存在缺失，直接延缓了安全事件的响应与补救效率；同时，用户信任度持续下滑，尤其在合法厂商的置换域名被滥用的情况下，将引发品牌与服务的信誉危机。

针对网安公司的社交工程攻击与试探

第三季观察到攻击者有意将目标瞄准「网安公司」或其公开服务窗口，常见攻击路径与手法可区分为两大类：

1. 长期潜伏式 Web / 服务窗口渗透

攻击者尝试以恶意软件感染服务端，成功感染后，以小量、频繁的请求（通常透过 HTTPS / 443）取得后续恶意程序，其核心目标是建立可长期维持的后门或定期搜集目标主机信息，并定期将信息上传到特定外部站点。攻击行为刻意低调（低频率、分散来源 IP、混淆 User-Agent），以避免被实时侦测系统标为异常流量。

2. 社交工程与商务洽谈伪装

以「购买服务」、「产品咨询」或「技术合作」之名接触业务承办人，诱导其提供企业内部信息、技术细节或测试访问权限。手段常结合精致的话术、仿真的公司文件与伪造联络人信息，单靠表面鉴别难以立即识破。

常见破绽

- 发信来源与真实性不一致

不少攻击使用的邮件并非来自他们声称的公司域名或官方邮件流程,若针对邮件头、来源 IP 与 SPF / DKIM / DMARC 进行核查,仍可发现破绽。

- 窗口响应机制缺少验证

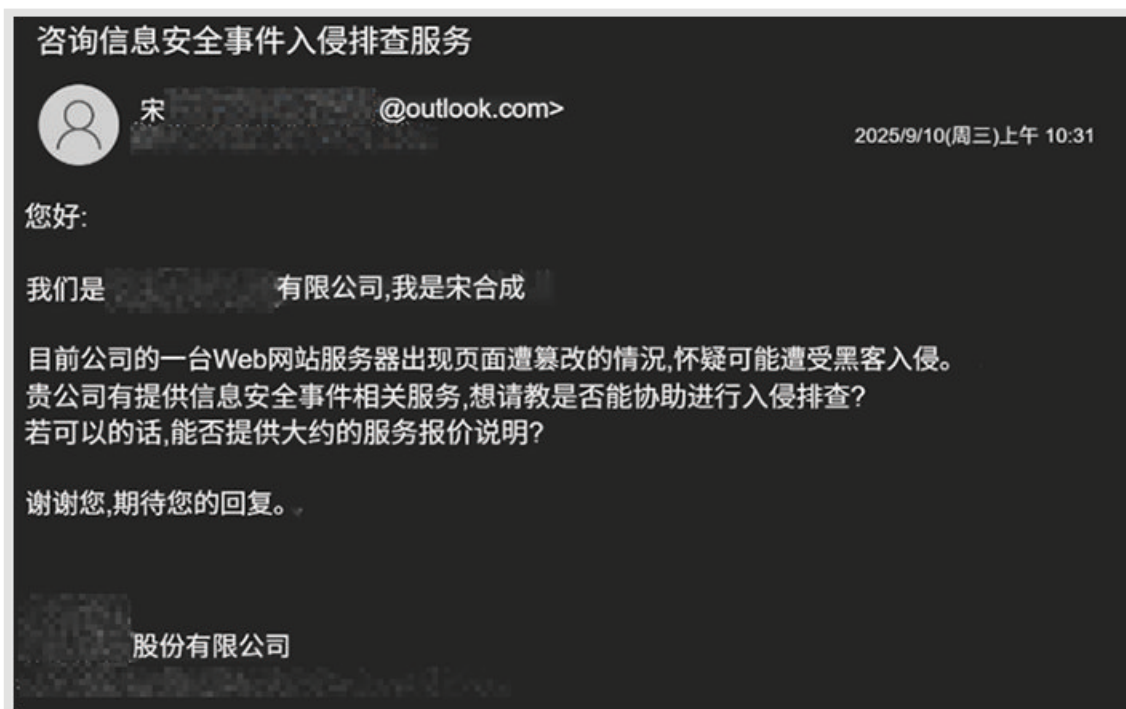
公司若以网页窗口作为第一接触点,但未对回复者进行强制验证(例如电话回拨、企业邮件网域验证或商业凭证),将提高被社交工程骗取信息的风险。

- 内部信息过度披露

公开的 FAQ、技术支持说明或产品文件若包含过多架构或技术细节,能被攻击者快速收集并用于定向攻击。



▮ 试图在服务窗口的计算机上埋入可以长期潜伏并泄资的后门



▣ 试图诱骗相关业务承办人泄露过多的讯息或技术情报

AI 进化带来的威胁

AI 与大型语言模型 (LLM) 在 2025 年下半年已广泛被攻防双方采用,对电子邮件安全的影响主要有三个方面:

1. 强化社交工程内容产出

AI 可生成高质量、针对特定组织或个人语气与文化语境的邮件文案,包含合理的时间脉络、专业术语与称谓,有效提高钓鱼与伪装成功率。并且能自动化 A / B 测试邮件标题、内容。诱使被害人进一步执行操作,以此快速迭代优化,以提高欺骗率。

2. 自动化漏洞发现与攻击链组合

攻击者利用 AI 加速漏洞扫描、解析邮件服务器或附件的潜在弱点,并自动生成对应利用代码或 payload。当 AI 结合自动化工具(如脚本、代理、多级跳转生成器)时,可以大规模产生变异化攻击,使传统签名方式防御失效。

3. 对企业内部 AI 系统的滥用(prompt injection / 隐藏指令)

随着企业导入 AI 助手处理邮件(如自动摘要、回复建议、敏感信息检测),攻击者可能在邮件正文中嵌入隐蔽指令(例如极小字体、白底白字、或特殊格式),诱使 AI 揭露敏感信息或执行不当行为(称为 prompt injection)。若 AI 的输出未经适当的审核或上下文限制,可能成为内部数据外泄或错误自动化决策的来源。

电子邮件攻防迈入新阶段

电子邮件攻防正进入「以“合法性”为盾、“自动化”为矛」的新阶段：攻击者大量利用 AI 生成社交工程与自动化工具，并利用合法第三方与置换链接的信任层来掩护恶意路径；同时，网安供应链本身与对外窗口成为高价值目标。单一层级的静态防御（例如只靠 SPF / DKIM / 传统过滤）已不足以应付这类复合、动态攻击。

未来趋势预测

- 多级跳转与合法服务滥用将更加普遍，防护会从「域名信誉」转向「跳转链分析」与「行为得分」。
- 攻击者对网安企业与对外窗口的试探会持续，促使网安公司本身采用更多“对抗式”自我测试与对外服务窗口强化（hardening）。
- AI 相关的 prompt injection 与模型滥用将成主要攻击向量，企业若不设限，AI 反而可能成为数据泄露的帮凶。

企业防护建议

- 强化身份与接触验证流程 - 对外业务 / 客服 / 窗口响应，应采用多因子验证与实体回拨核实。
- AI 使用原则与防护 - 针对内部 AI 处理邮件的流程。设计输入净化、输出审核与最小授权。
- 钓鱼演练与社交工程防御训练 - 针对 VIP / 财务 / 客服进行定向演练与应急响应培训。
- 建立业界协作与实时通报机制 - 当发现被滥用的第三方或置换域名，应快速通报、分享入侵指标 IoCs 与同步封锁。

关于守内安

守内安信息科技(上海)有限公司,是上海市政府及国家奖励支持的自主研发高科技创新的“双软认定企业”和“高新技术企业”,钻研邮件风险管理和信息安全内控管理。以电子邮件安全管理为核心,研发了一系列“电邮安全与合规”为中心的核心产品线,衍生到威胁防御与联合防御体系。守内安 20 年来秉承“以客为尊”的服务理念,树立了“服务·品质·值得信赖”的品牌理念,目前已拥有 7000+ 家全球性企业级用户,终端用户达 80,000,000+ 人次。

