

# Qualys TotalAI

降低 Gen AI 和 LLM 工作负载的风险

“在AI时代，最大的风险不是不去创新，而是在没有稳固基础的情况下创新。”

生成式人工智能（Gen AI）和大语言模型（LLM）正在革新各行各业，但是，它们的快速应用带来了严峻的网络安全挑战，为陈旧系统构建的传统安全方法根本无法应对。如今，企业面临着知识产权被盗、数据泄露、违反隐私法规等风险。在这种情况下，就更需要了解 LLM 的所在位置、漏洞以及暴露程度。这正是 **Qualys TotalAI** 发挥作用的地方。



**Qualys TotalAI** 为企业提供针对 AI 生态系统的可视化和控制，可保障企业运营安全并减轻风险。它旨在应对特定威胁，例如，**提示注入（Prompt Injection）、越狱（Jailbreaks）、模型被盗、敏感信息泄露**，让 AI 基础设施保持稳健且合规。

**Qualys TotalAI** 为 AI 工作负载提供全面的保护，可应对 **OWASP** 列出的 **十大 LLM 应用程序风险**。



## AI 工作负载发现与监察

通过 AI 资产映射和创建 AI 资产清单，实现可视化管理。全面了解所有 AI 工作负载，包括可能带来隐藏风险的影子模型（Shadow Models）。



## AI 漏洞管理

区别于传统扫描方式，**Qualys TotalAI** 使用 650+ 种 AI 特定检测，评估 AI 模型和基础设施，防止数据被盗、配置错误和其他风险。



## LLM 扫描

评估 LLM 模型是否存在提示注入、越狱、模型被盗等重大风险。防范可能危害数据或知识产权的威胁。



## 使用 TruRisk™ 确定优先级并做出响应

利用 TruRisk™ 评分工具，集中精力处理重要事项。优化修复工作，让最严重的风险得到优先解决。



## 合法合规

确保 AI 模型不违反 GDPR 和 CCPA 等数据保护法规，防止数据泄露，避免高额罚款。