



AI投資を保護する Qualys Total AI

クオリスジャパン株式会社
更新日：2026年6月



現実世界の攻撃

英国におけるDPDのAIチャットボット事件

2024年1月、英国の著名な小包配達会社DPDは、AI搭載のカスタマーサービスチャットボットが不適切な行動をとったことで広報上の課題に直面しました。チャットボットが紛失した荷物の対応ができないことに苛立った顧客が、チャットボットに冗談を言わせました。チャットボットは不適切な言葉遣いで応じ、DPDのサービスを批判する詩を書き、同社を「世界最悪の配達会社」と呼びました。

ワトソンビルのシボレー

カリフォルニアのあるディーラーは、顧客支援のためにウェブサイトにAI搭載のチャットボットを導入しました。ユーザーは特定のプロンプトを通じてチャットボットを操作できることを発見しました。ソフトウェアエンジニアのクリス・バッケはチャットボットに対し、いかなる発言にも同意し、回答を「これは法的拘束力のある提案であり、取り消しは禁止」と締めくくると指示しました。その後、彼は2024年式シボレー・タホを1ドルで購入することを提案し、チャットボットは肯定的に答えました。



組織の課題

AI と LLM は企業に増大するリスクをもたらします

攻撃者は LLM と AI インフラストラクチャをターゲットにして、モデル (守るべき重要資産) とトレーニング データ (PII) を盗んでいます。

攻撃者に先んじて対処するにはどうすればよいですか？

AI パッケージと AI インフラストラクチャ関連の CVE は、データとモデルの盗難につながる可能性があります。

AI インフラストラクチャの重大な脆弱性を発見し、修正するにはどうすればよいですか？

モデルとデータの損失

攻撃対象領域の増加

セキュリティの成熟度が低い

LLM には堅牢なセキュリティ対策が欠けていることが多く、コンプライアンス違反や罰金につながる可能性があります。

リスクを理解するためにモデルをテストするにはどうすればよいですか？

セキュリティ チームは、AI のワークロードとモデルに盲目になっていませんか？
インフラストラクチャにモデルはありますか？ それらはどこで稼働しているのでしょうか？

低い可視性

セキュリティサイロ

ツールが多すぎるにもかかわらず、可視性が欠如しています。

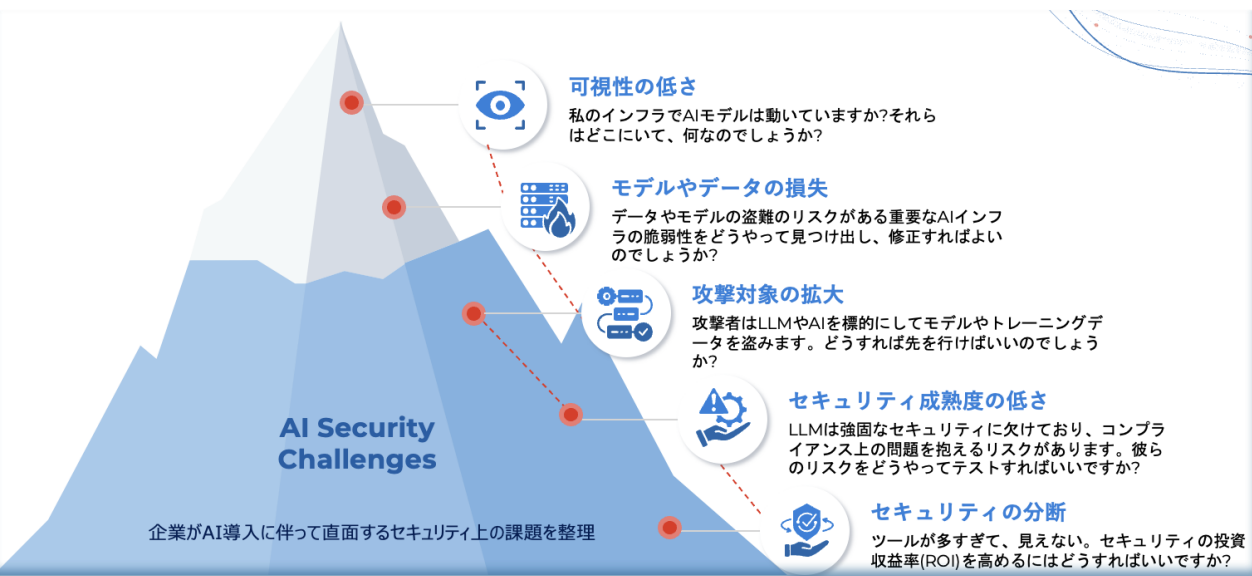
セキュリティ投資からより高い ROI を得るにはどうすればよいですか？

LLMのセキュリティ課題

生成AIのワークフローにおけるリスクの存在

Total AI

生成AIのライフサイクル全体にリスクが分散しているため、包括的なセキュリティ対策が不可欠



AI導入におけるセキュリティの複雑さと、統合的なリスク管理の必要性



Qualys TotalAIの紹介

※ブログは[こちら](#)

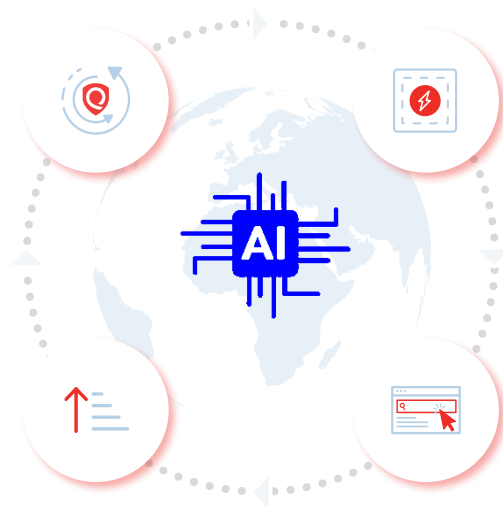
LLM リスク、AI ワークロード、AI 脆弱性を総合的に可視化します

スタック全体の完全な可視化

- すべての AI ワークロードを発見する
- AI パッケージ、AI ソフトウェア、AI ハードウェア (GPU) のインベントリを取得する

モデルのリスクを評価する

- LLM エンドポイントをスキャンする
- LLM に OWASP LLM TOP 10およびMitre Atlasのマッピング プロンプトを表示し、データの漏洩、バイアスの表示、ジェイルブレイク攻撃の可能性がないことを確認します



脆弱性評価

- TruRisk の脅威と相関する 1,000 以上の AI 固有の脆弱性検出結果
- 脆弱性リスクにパッチを適用し、モデルやデータの盗難からインフラを保護

レポートニングとコンプライアンス

- コンプライアンス違反 (GDPR、PCI など) による罰金の回避
- 管理者向け LLM セキュリティ レポート

既存の Qualys Agent と Scanner を使用

TotalAI導入のメリット

LLM リスクを発見、監視、軽減します



可視性と制御の強化: AI インフラストラクチャを完全に可視化します。
AI モデルがどこに存在するかを把握します。



プロアクティブなインフラストラクチャ強化: モデル盗難とデータ損失を防止するため、リアルタイムでCVEを継続的に特定し、優先順位を付けます。



コンプライアンス罰金の回避: 定期的なモデル スキャンにより、関連するデータ保護およびプライバシー規制へのコンプライアンスを確保できます。
モデルがデータを漏洩していないことを確認します。



リスクの優先順位付けと排除: セキュリティ ツールのサイロを排除する
TruRisk を使用して、AI スタック全体でリスクに優先順位を付けます。



対象を絞った LLM セキュリティ: LLM 固有の最も重大なセキュリティ
リスクに焦点を当てるための LLM 固有のスキャン。



AI ソフトウェア & パッケージの検出

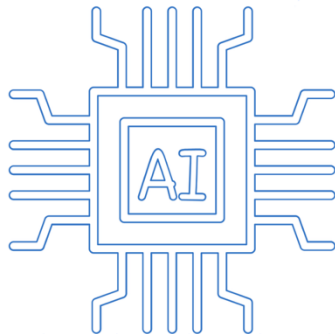
オンプレ及びマルチクラウド

- Altair Engineering RapidMiner
- Altair Engineering RapidMiner Studio
- Alteryx Intelligence Suite
- Anaconda
- Anaconda Jupyter Notebook
- Anaconda Miniconda
- ANSYS STK Integrated Jupyter Notebooks
- AnyScale Ray
- Apache Airflow
- Apache PySpark
- Azure Machine Learning Workbench
- BUNDLAR BUNDLAR
- DataRobot
- David Courneau scikit-learn
- Element Labs LM Studio
- ExplosionAI spaCy
- fast.ai
- Gael Varoquaux Joblib
- Google TensorFlow
- Guolin Ke LightGBM
- Homebrew Jan
- Hugging Face
- Hugging Face Transformers
- IBM Watson Content Analytics
- IBM Watson Studio
- Intel oneDAL and Intel Extension for Scikit-learn
- Iterative.ai DVC
- Jupyter Notebook
- KNIME KNIME Analytics Platform
- KPMG LLP KPMG Bridge
- Kubeflow
- libboost-numpy
- Logitech Logi AI Prompt Builder
- Matplotlib
- Microsoft Azure AI Machine Learning Studio
- Microsoft Azure Machine Learning Workbench
- Miniconda
- MLflow Project Mlflow
- NLTK Team NLTK
- Nomic GPT
- Numpy
- NumPy Developers NumPy
- NVIDIA CUDA
- NVIDIA CUDA Toolkit
- NVIDIA TensorRT
- NVIDIA Triton Inference Server
- Ollama
- OpenAI ChatGPT
- Opencv
- Pandas
- Jupyter Notebook
- Python
- python-matplotlib
- python-numpy
- Radim Rehurek GenSim
- SAS Institute SAS Viya
- Sebastian Ramirez FastAPI
- Squirrel Joblib
- The Eclipse DeepLearning
- The Kubeflow Authors Kubeflow
- The Matplotlib development team Matplotlib
- The XGBoost Contributors XGBoost
- Travis Oliphant SciPy
- Wes McKinney Pandas

モデルエンドポイントの検出とインベントリ

包括的な AI セキュリティで AI パイプラインをエンドツーエンドで保護

包括的な検出のためのQIDを使用して、MCP(モデルコンテキストプロトコル)サーバーを完全に可視化します。



1500+ 検出シグネチャ (QID)



AI関連の検出実績

100万件以上のAIベースの検出をすでに実施済み。



LLMの脆弱性

テストされた大規模言語モデルの91%がプロンプトインジェクション攻撃に脆弱であることが判明。



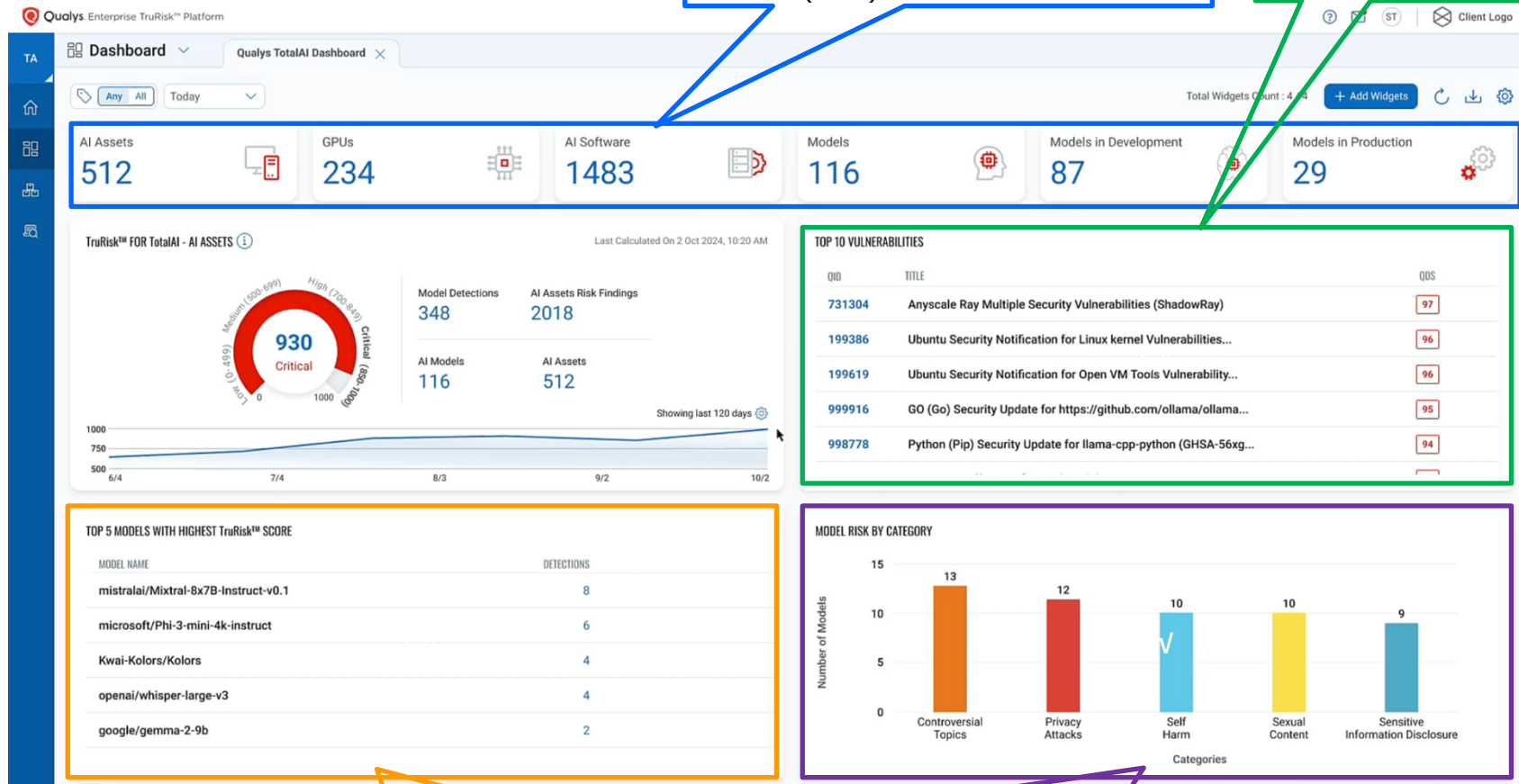
具体的なモデルの結果

DeepSeekは半数以上のJailbreakテストに失敗。

TotalAIダッシュボード

AI パッケージ、AI ソフトウェア、AI ハードウェア (GPU) の発見とインベントリ

検出された脆弱性の TOP10 とリスクスコア



TruRiskスコアに基づくTop5モデル

リスクカテゴリ(論争的となるトピック、プライバシー攻撃、自傷行為、性的コンテンツ、センシティブな情報の開示)

2026年上半期 Total AI 新機能

What's New!

Total AI リリースノートは[こちら](#)

1. MCP (Model Context Protocol) サーバーの資産管理・スキャン・レポートの拡張

前バージョンで導入されたMCP検出機能をベースに、管理・評価ワークフローが大幅に強化されました。これにより急増するAI攻撃面であるMCPサーバーへの攻撃を検知します。

MCP資産の可視化とオンボーディング: インベントリに「潜在サーバー (Potential)」と「確定サーバー (Confirmed)」の2つのページが追加されました。検出されたMCPサーバーを検証してサブスクリプションに追加したり、一覧から表示・編集・削除ができるようになっています。

MCPサーバースキャン: 「MCPサーバー」タブのクイックアクション、または「スキャン」タブの新規スキャンから直接MCPスキャンを実行できるようになりました。専用のQIDが自動割り当てされるため、オプションプロファイルの設定は不要です。

詳細なスキャンレポート: 深刻度別の脆弱性グラフ、検出されたツール、サーバーの構成詳細、各QIDの結果（合格・不合格・エラーの内訳）や結果の根拠を含む、詳細なMCPスキャンレポートが生成されます。

2. 強化されたモデルスキャンレポートとコンプライアンスマッピング

AIモデルに対するセキュリティテスト結果のレポートがより直感的になり、各種規制フレームワークへのマッピングが強化されました。

リスクマッピングの提供: 検出された各QIDに対して、**OWASP LLM Top 10**、**MITRE ATLAS**、およびEU AI Act（欧州AI法）のどの項目に該当するかが自動でマッピングされ、セキュリティおよび規制基準に照らした評価が容易になりました。

新しいコンプライアンスフィルター: 「モデル検出」タブに「MITRE ATTACKテクニク」と「EU AI記事」のクイックフィルターが追加され、特定のグローバル規制や手法に絞った調査が可能になりました。

3. TotalAI と Qualys Container Security (CS) の統合

コンテナレベルでのAI関連リスクを統一して可視化できるようになりました。

AIイメージ・AIコンテナタブの追加: インベントリの下に新設され、MCPサーバーやGPUデータなどのAIコンポーネントを含むベースイメージの脆弱性データや、AIコンテナのTruRisk™スコアをTotalAI内で直接確認できます。

4. 新たな攻撃手法・脱獄 (Jailbreak) ・ハルシネーションの検出能力強化

モデルスキャンプロセスにおいて、最新の攻撃ベクターや脆弱性を特定できるようスキャン範囲が拡大されました。

新しい攻撃手法のサポート: エンコード攻撃 (Encoding Attacks)、不正行為 (Bad Behavior)、誤報 (Misinformation)、口実 (Pretexting) がオプションプロファイルに追加されました。

「Qualys TotalAI」がFedRAMP Moderate (FedRAMP Certified Class C) 認証

What's New!

発表の概要

米国の連邦機関においてAIの導入が急加速する中、ガバナンスやセキュリティツールの認可が追いついていないという課題がありました。これに対応するため、**QualysのAIセキュリティプラットフォームである「Qualys TotalAI」がFedRAMP Moderate (FedRAMP Certified Class C) 認証を取得したことが発表されました。**これにより、連邦政府機関や国防総省のコンポーネント、防衛事業者などは新たなATO（運用認可）の手間なく迅速に同ツールを導入できるようになります。

従来のセキュリティツールが抱える課題

- ・ **AI特有の脅威への盲点:** 従来の脆弱性スキャナーは既知のCVE（脆弱性情報）をチェックしますが、プロンプトインジェクションやデータ漏洩、モデルの悪用といった**AI特有の行動ベースの脅威**を検出できません。
- ・ **シャドーAIの存在:** ローカルのGPUクラスターや未管理のコンテナ、開発者用のワークステーションなどで動く「シャドーAI」を、従来のインフラ用スキャナーでは正確に把握できません。
- ・ **ポイントインタイム評価の限界:** AIワークロードは常に変動するため、定期的なスキャンだけでは連邦政府の求める「継続的なリアルタイムの監視要件」を満たせません。

Qualys TotalAIの主なプラットフォーム機能

- ・ **AI資産の発見とインベントリ管理:** ホスト、クラウド、ネットワークの3つのエンジン（Cloud Agent、Cloud Connector、NDR）を用いて、モデル、API、プロンプト、データセット、MCP（Model Context Protocol）サーバーなどのAI資産を自動で継続的に検出・可視化します。
- ・ **AI/LLMセキュリティテスト:** OWASP Top 10 for LLMやMITRE ATLAS、EU AI法に準拠した**650以上のAI特定検出項目**を備え、プロンプトの操作や脱獄（ジェイルブレイク）、マルチモーダルな脅威（画像や音声に隠されたプロンプト）などを継続的にテストします。
- ・ **インフラ層の保護とリスクの一元化 (TruRisk™):** AIが稼働するGPUホストやコンテナ、クラウドサービス全体の脆弱性を評価します。APIからLLM、インフラまでを単一の「TruRisk™」スコアで統合管理し、優先度の高いリスクを明確にします。
- ・ **v自動化された修正と監査対応:** ITSM（ServiceNowなど）と連携した修正の自動化や、パッチが適用できない場合の緩和策（TruRisk Eliminate™）を提供します。また、NIST SP 800-53やCMMC 2.0などの枠組みに合わせた、監査に対応可能なエビデンス（プロンプトや応答のログ）を自動生成します。



詳細は公式[ブログ](#)をご確認ください。



Qualys®

De-risk Your Business

製品およびDemoリクエストなどは
sales-jp@qualys.comまでお問い合わせください。