

AI の安全性を証明できますか？



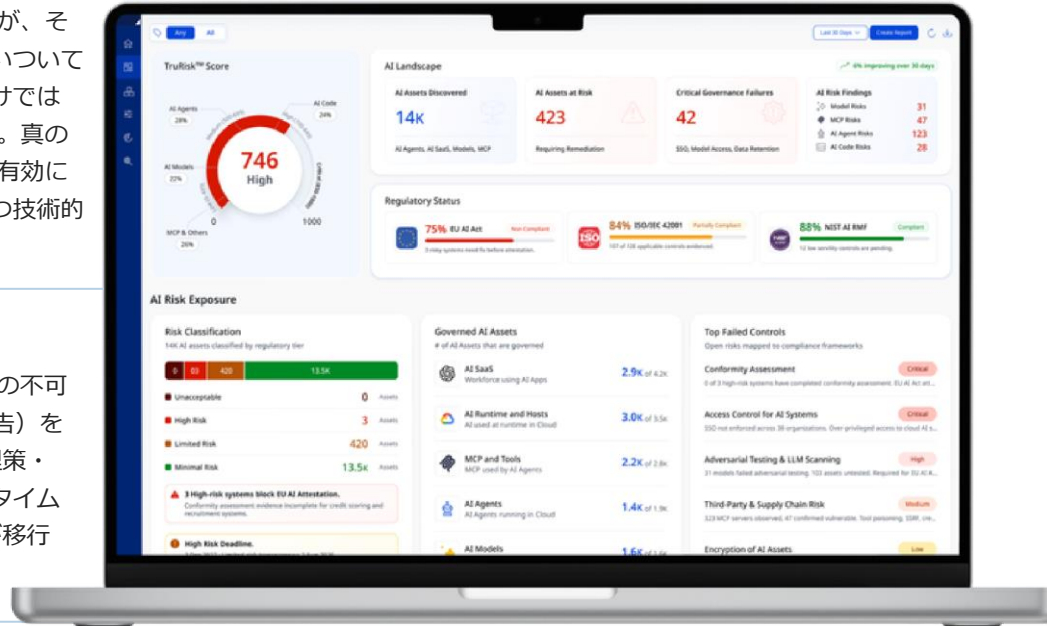
AI イノベーションを拡大し、安全性を確保する。

AI ガバナンスが機能していることを証明しましょう。

AI は今や企業のインフラとなっていますが、その導入スピードにセキュリティ対策が追いついていないのが現状です。静的なポリシーだけでは AI の安全性を証明することはできません。真の AI ガバナンスには、セキュリティ統制が有効に機能していることを裏付ける、継続的かつ技術的な実証が不可欠です。

## Qualys TotalAI

Qualys TotalAI は、AI ガバナンス運用の不可欠な 4 つの要素（発見、評価、修復、報告）を統合することで、断片化された AI の管理策・ポリシー・ツールから、コードからランタイムに至る包括的な AI リスク管理へと組織が移行できるよう支援します。



960+ AI 特化型検知

300+ 攻撃シナリオ

55+ ポスチャ制御

### 目に見えない AI を可視化する

- ホスト、クラウド、MCP、コンテナ、パイプライン、SaaS にわたる、承認済み AI およびシャドーAI のインベントリを継続的に構築します。
- すべての資産を、所有者、運用コンテキスト、エクスポージャー、リスクと関連付けてマッピングします。
- GPU、モデル、コンテナ、MCP サーバー全体で、Qualys の AI フィンガープリント技術を活用します。

### 最も重要な AI リスクへの対処

- 重大度、モデルの挙動、エクスポージャー、およびランタイム時の証拠を関連付け、有害なリスクの組み合わせを特定します。
- 実行可能なガイダンスと、即座に導入可能なガードレールを、担当者に直接提供します。
- 修正内容を再テストし、リスクが解消されたことを示す、説得力のある証拠を記録します。

### あらゆる段階で AI のセキュリティリスクを管理

- **ポスチャとコード:** 設定ミス、過剰な権限、シークレット、脆弱性のある依存関係を特定します。
- **挙動:** LLM や MCP サーバーに対し、プロンプトインジェクション、ジェイルブレイク、ツールポイズニングのテストを行います。
- **ランタイム:** カーネルレベルの eBPF や、ブラウザおよびコネクタからの SaaS テレメトリを活用し、エージェントの実行状況やデータの移動を監視します。

### エビデンスに基づく AI ガバナンスの報告

- 継続的に収集されるエビデンスを、取締役会、顧客、規制当局への報告に適した(監査対応可能な)形式に変換します。
- インベントリ、テストでの不適合、リスクの優先順位付け、および再テストによる是正完了状況を文書化します。
- 検出された事項を、NIST AI RMF、EU AI 法、OWASP Top 10 for LLMs、MITRE ATLAS といった枠組みと関連付けます。

FedRAMP Moderate Authorized

TotalAI は現在、FedRAMP Moderate の認定を受けて導入可能であり、NIST AI RMF、NIST SP 800-53、および CMMC 2.0 に対応した監査可能なエビデンスを備えています。



今すぐトライ

自社の AI フットプリントを把握し、AI の攻撃対象領域を保護しましょう。  
30 日間で AI ガバナンスの基盤を構築します。