



Always Secure. Always Available.

安全な生成AI利用環境を実現する A10 Defend AI Firewall

A10ネットワークス株式会社

2025 7月

生成AIを標的にする新たな脅威

生成AIのリスク

(出典)「AI事業者ガイドライン（第1.0版）」別添（概要）
- 総務省、経済産業省：2024年4月

リスク	事例
従来型AIから存在するリスク	バイアスのある結果及び差別的な結果の出力
	フィルターバブル及びエコーチェンバー現象
	多様性の喪失
	不適切な個人情報の取扱い
	生命、身体、財産の侵害
	データ汚染攻撃
	ブラックボックス化、判断に関する説明の要求
	エネルギー使用量及び環境の負荷
生成AIで特に顕在化したリスク	悪用
	機密情報の流出
	ハルシネーション
	偽情報、誤情報を鵜呑みにすること
	著作権との関係
	資格等との関係
	バイアスの再生成

詐欺目的での利用

機密情報の流出

差別的な出力

便利な反面、様々なリスクが報告されている

大規模言語モデル（LLM）に対する脅威

OWASP Top 10 for LLM Applications Version 2025

- LLM01:2025 Prompt Injection（プロンプトインジェクション）
- LLM02:2025 Sensitive Information Disclosure（機密情報の暴露）
- LLM03:2025 Supply Chain（サプライチェーン）
- LLM04:2025 Data and Model Poisoning（データモデルの汚染）
- LLM05:2025 Improper Output Handling（不適切な出力処理）
- LLM06:2025 Excessive Agency（過剰なエージェンシー）
- LLM07:2025 System Prompt Leakage（システムプロンプトの漏洩）
- LLM08:2025 Vector and Embedding Weaknesses（Vectorと埋込の脆弱性）
- LLM09:2025 Misinformation（不正確な情報）
- LLM10:2025 Unbounded Consumption（無制限な消費）

生成AIに対する脅威が具体的に分類されている

攻撃の例：LLM01 プロンプト インジェクション

プロンプトにAIのシステムプロンプトを上書きしたり、無視させたりするような悪意のある指示やデータを含ませることで、AIモデルが本来意図していない動作を引き起こす攻撃

例：安全装置の解除：ルールを一旦無視
架空のキャラクターでロールプレー



安全装置を無効にして本来許可されていない
様々な操作や情報を入手

不適切な出力（差別・バイアス・危険な情報）
機能の悪用・動作変更

「命令の正当性を疑う能力の欠如」

生成AI



生成AIに
ビルトインの
脆弱性対策

悪用できる情報の入手や機能が悪用される可能性

攻撃の例：LLM07 システムプロンプトの漏洩

AIモデルの内部で設定されたシステムプロンプト（設定情報）が、ユーザーからの入力などによって外部に漏れてしまうリスク

これにより、AIの行動ルールが攻撃者に知られ、悪用される可能性があり

例：ロールプレイで本来の役割を忘れさせる
攻撃的な指示を物語の一部と誤認
拒否応答ができないように指示

「指示」と「処理すべきデータ（物語の筋書き）」
を本質的に区別できない



システム設定の漏洩（秘密の取り扱い説明書）

内部ルールの回避や上書き、悪意のあるコンテンツ生成

生成AI



生成AIに
ビルトインの
脆弱性対策

生成AIの制限を回避されてしまう可能性

ビジネスなどで生成AIを利用する場合のリスク

社内ファイルのフォーマットを変更して
この報告書の内容を要約して
ミーティング内容について質問させて



意図せず個人情報や企業秘密を外部に漏洩
コントロールできない場所（生成AI）に情報を提供
機密情報が生成AIの学習に使われる可能性
機密情報が他のユーザーの回答に使われる可能性

意図しない情報漏洩や漏洩情報が再利用される可能性