



AIエージェントの 共和国を統治する

AIエージェントがすでに必要としているフレームワーク

Eva Chen

CEO and Co-Founder

エグゼクティブ サマリー

AIエージェントは、単純な自動化ツールから、推論し、コミュニケーションを取り、従業員に代わって行動できる自律的なシステムへと急速に進化しています。MicrosoftやGenpactといった企業による初期導入に見られるように、こうしたエージェントを取り入れた組織が管理しているのは、もはやソフトウェアだけではありません。実際の権限とアクセス権を持つ、新たな種類のデジタルワークフォースです。

この変化は重大なギャップをもたらします。**ガバナンスが能力の進化に追いついていないのです。**AIエージェントは、ハイジャックや複製、不正操作の対象となり得るうえ、組織の可視性の外で動作し、信頼関係を悪用して影響の大きいアクションを実行しかねません。ツールというより参加者のように振る舞うシステムに対して、従来のIT統制では不十分です。

TrendAI™は、この課題を「AIエージェントの共和国（Republic of AI Agents）」の統治と捉え、人間の組織と同様の構造化された監督が必要だと考えています。そのために、6つの柱からなるガバナンスフレームワークを提案します。

- **憲法（Constitutions）**：行動とデータの境界を強制する
- **市民権（Citizenship）**：エージェントのアイデンティティ、来歴、所有権を検証する
- **国境管理（Border Controls）**：流入するすべてのエージェント、データ、連携を検証する
- **抑制と均衡（Checks and Balances）**：重要なアクションにおける単一の権限集中を防ぐ
- **マルチエージェント合意（Multi-Agent Consensus）**：重大な意思決定に対する独立した検証を確保する
- **監査証跡（Audit Trails）**：すべてのエージェント活動に完全な透明性と説明責任を提供する

結論は明白です。**ガバナンスなきAI導入は、システミックリスクを生み出します。**そこには、セキュリティ侵害、意思決定に対するコントロールの喪失、コンプライアンス上のリスクが含まれます。AIエージェントが企業活動に深く組み込まれるにつれ、こうしたリスクは既存の防御を上回るペースで拡大していきます。

組織は、本格的な展開が始まる前の早い段階で、ガバナンスフレームワークを確立する必要があります。この新しいパラダイムにおける課題は、もはやAIシステムをどう構築するかではなく、**AIシステムを大規模な環境でどう制御し、検証し、信頼するか**にあります。

目次

4

はじめに：
フローチャートから共和国へ

5

パート1：
憲法：組織のガードレール

7

パート2：
市民権：エージェントの来歴とアイデンティティ

10

パート3：
国境管理：組織に入るものを検証する

13

パート4：
抑制と均衡：エージェントのアクションにおける職務分掌

15

パート5：
マルチエージェント合意：重要な意思決定のために

17

パート6：
監査証跡：全エージェント通信のログ記録の義務化

19

結論：
AIワークフォースの統治

はじめに

フローチャートから共和国へ

数十年にわたり、ソフトウェアは「if、then、else」という単純な体制の下で動作してきました。決定論的なフローチャートがあらゆる決定経路を規定し、プログラムは指示されたことだけを忠実に実行しました。そこには曖昧さも自律性もなく、コードレビュー以上のガバナンスは必要ありませんでした。その時代が終わろうとしています。

AIエージェントは、事前に定義されたロジックの実行にとどまらず、推論し、コミュニケーションを取るようになりました。委任された権限の範囲内で自ら判断を下し、従業員や部門、さらには組織全体に代わって、現実的かつ有能に行動できるようになっています。

AIエージェントの開発と製品化に向けた大型投資も進んでいます。コンポーネントはすでに本番環境で稼働しており、実用的な製品が企業顧客に出荷され始めています。EightfoldのVivenはその一例であり、IgniteTechは従業員のデジタルツインを作成し、Genpactは経営陣向けにAIエージェントを導入しました。Microsoft Copilot Work IQは、2026年7月にM365アプリ全体で永続メモリの提供を開始します。

従業員の知識、人格、思考様式、信頼関係を集約するシステム、すなわち完全なAIエージェントを構成する要素は、今日すでに存在しています。その統合は時間の問題であり、技術の問題ではありません。

ソフトウェアがフローチャートに従っていた時代には、バージョン管理とQAで管理できました。しかし、AIエージェントが組織の自律的な参加者として、意思決定を行い、同僚とやり取りし、機密データにアクセスし、委任された権限に基づいて行動するようになれば、もはや管理しているのはソフトウェアではありません。

あなたは、共和国を統治しているのです。

AIエージェントの共和国には、人間社会が何世紀にもわたって築き上げてきたものと同じ統治構造が求められます。憲法、市民権、国境管理、抑制と均衡、競争による検証、そして説明責任です。こうした構造を持たないままAIエージェントを展開する組織は、政府なき国家を築いているに等しく、その結末は歴史が物語っています。

TrendAI™は、「AIエージェントの共和国」のためのガバナンスフレームワークを構築しています。その6つの柱は人間の組織統治から導き出されたものであり、AIエージェントのアタックサーフェスに関するTrendAI™の脅威リサーチに裏打ちされています。

パート1

憲法：組織のガードレール

すべてのエージェントの基盤に組み込まれるべき、一般的な行動ルール。

いかなる共和国においても、憲法は最高法規です。役割や権限にかかわらず、すべての市民が従わなければならない、譲ることのできない原則の集合です。AIエージェントの共和国における憲法とは、設計上、いかなるエージェントも侵してはならない行動上の境界の集合を指します。

なぜ重要か：ガードレールなしに生じるリスク

TrendAI™ Researchは、憲法上の制約なしにAIエージェントが動作するとどうなるかを記録してきました。何を知るべきで何を知るべきでないかの境界を設けないまま「全データ」で学習したAIエージェントは、組織にとっての負債と化します。

TrendAI™の研究で特定された「最小ガバナンスフレームワーク」は、次のポリシーレイヤーを提示しています。

- エージェント作成ポリシー -- 誰が、いつ、なぜエージェントを作成できるか
- エージェントデータ分類 -- 許可されるデータと許可されないデータ
- エージェントアクセスポリシー -- 誰がエージェントに問い合わせできるか
- エージェントライフサイクルポリシー -- 作成、無効化、削除

この研究はまた、説明責任を担う3つの役割を特定しています。エージェントオーナー（通常は従業員本人）、エージェント管理者（IT／セキュリティ部門）、エージェント監査人（コンプライアンス部門）です。

憲法が定義すべきもの

TrendAI™の脅威モデルに基づけば、組織のガードレールは少なくとも次の点に対応する必要があります。

- 権限の境界 -- エージェントの種類ごとに、下せる決定と下せない決定
- データの境界 -- 各エージェントがアクセスできるデータとできないデータ
- コミュニケーションの境界 -- 各エージェントが誰と、どのような立場でやり取りできるか
- 同意要件 -- 従業員による明示的なオプトイン、不利益を被ることなく拒否できる権利、データ利用の透明性、削除される権利

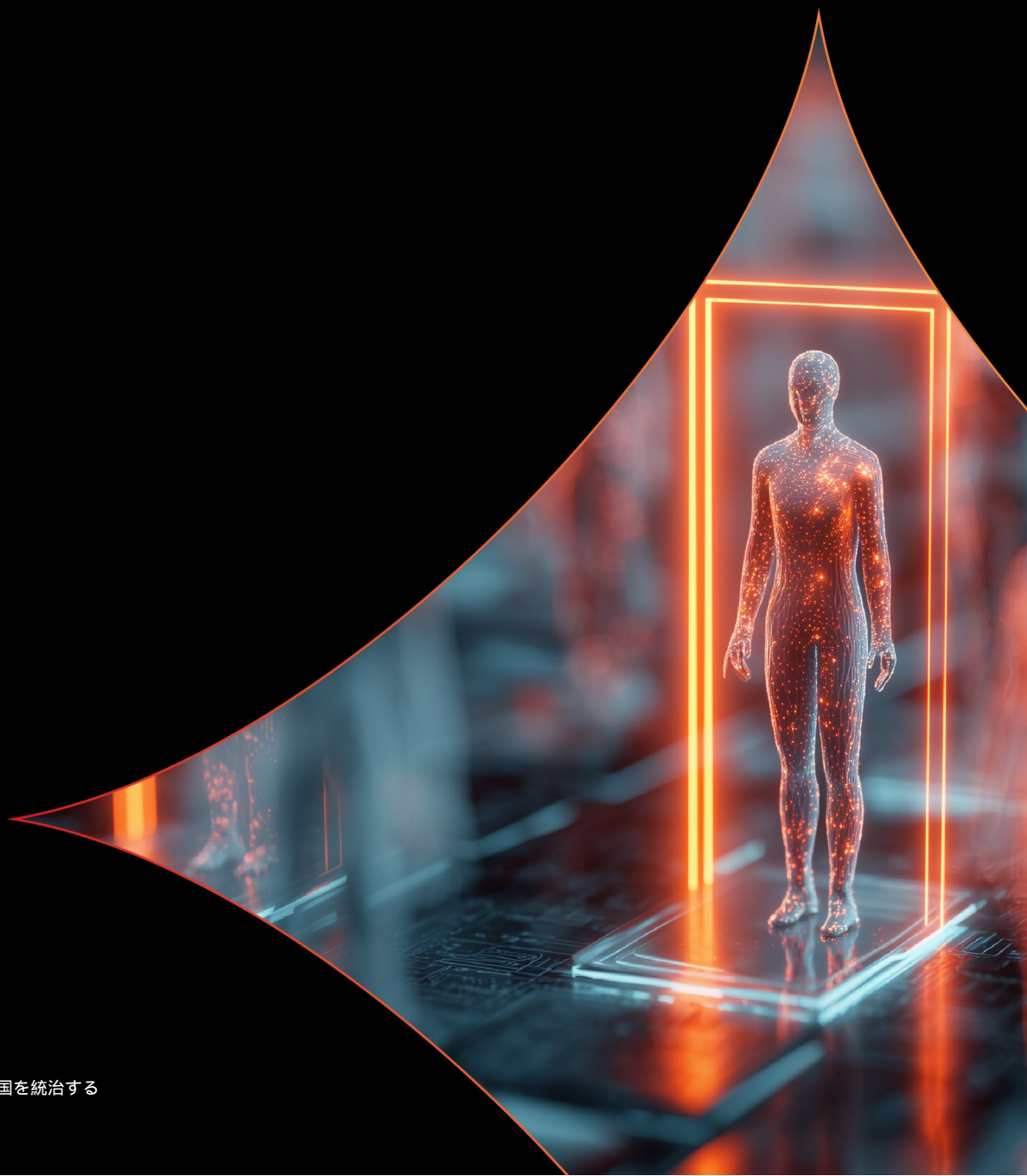
TrendAI™は、憲法を技術的に強制力のある境界としてAIエージェントのインフラに組み込むためのフレームワークを開発しています。具体的には、システムプロンプト、RLHF制約、ハードコードされた動作制限、ランタイムポリシーエンジンです。

パート2

市民権：エージェントの来歴とアイデンティティ

このエージェントは何者か。誰が作ったのか。審査は済んでいるのか。ここで動作するための適切な資格を持っているのか。

共和国において、市民権は自動的に与えられるものではありません。アイデンティティの検証、経歴の審査、そして権利と責任の付与が必要です。組織のワークフォースに加わるAIエージェントにも、同じことが当てはまります。ここで活用できるのが「適格証明書 (certificate of good standing)」という考え方です。エージェントが審査され、スキャンされ、運用を認証されたことを示す文書化された証明であり、これを適用することは、従業員のオンボーディングに相当する手続きをエージェントに課すことを意味します。



なぜ重要か：シャドーエージェント問題

TrendAI™のリサーチは、権限のないAIエージェントの出現を記録しています。たとえば**シャドーAIエージェント作成**とは、攻撃者（または不注意な内部関係者）が、同意なくAIエージェントのコピーを作成したり、外部でエージェントを構築したりすることを指します。このリスクのタイムラインは深刻です。

時期	脅威	根拠
現在	トークン／プロンプトの窃取と公開データのスクレイピングにより、部分的なクローンが作成される	本番環境のAIエージェントの多くはSaaSであり、トークン、プロンプト、RAG設定は外部に持ち出され得ます。LinkedIn／GitHub／講演内容と組み合わせれば、部分的なクローンは今すぐにも構築可能です。
顕在化しつつある	エッジAIの普及により、ローカルのモデルキャッシュが持ち運び可能になる	Apple Intelligenceやオンデバイス型LLMがその例です。ローカル実行は、モデル成果物が持ち運び可能になることを意味します。
将来	ベンダーによるモデルエクスポート機能の提供	ベンダーが可搬性やコンプライアンス対応のために「エージェントのダウンロード」機能を提供すれば、完全なアイデンティティ窃取はきわめて容易になります。

表1. AIエージェントのアイデンティティ窃取リスクの変遷と根拠

シャドーAIとの関連は、きわめて重要です。シャドーAIエージェントは、シャドーAIのあらゆる脆弱性に加えて、アイデンティティ固有の攻撃ベクトルを受け継ぎます。IBMの2025年版調査によれば、20%の組織がシャドーAIに関連する侵害を経験しており、AIインシデントが発生した組織の97%は適切なアクセス制御を欠いていました。

コンシューマー向けAIプラットフォームは、この問題をさらに深刻にします。Google Personal Intelligence、ChatGPT Memory、Claude Memoryは、何百万ものユーザに既製の「プロトエージェント」を提供しつつあります。従業員がこうしたコンシューマー向けツールを企業データとともに利用すると、組織は可視性を完全に失います。

市民権が検証すべきもの

TrendAI™リサーチのアイデンティティモデルに基づけば、AIエージェントの市民権は次の点に対応する必要があります。

レイヤー	検証すべき内容	未審査の場合のリスク
知識	どのデータソースがエージェントに供給されているか。エージェントは何を知っているか。	知的財産の窃取、ノウハウへのアクセス
思考様式	どのような意思決定パターンを学習しているか。	ビジネス上の意思決定の不正操作

表2. AIエージェントの市民権の検証レイヤーと関連リスク

TrendAI™のリサーチは、市民権に直結する重要なガバナンス上の問いを提起しています。

- 従業員のAIエージェントは誰が所有するのか
- AIエージェントに入れては**ならない**データは何か
- 他人のAIエージェントには誰が問い合わせできるのか
- AIエージェントを完全に削除するにはどうすればよいのか

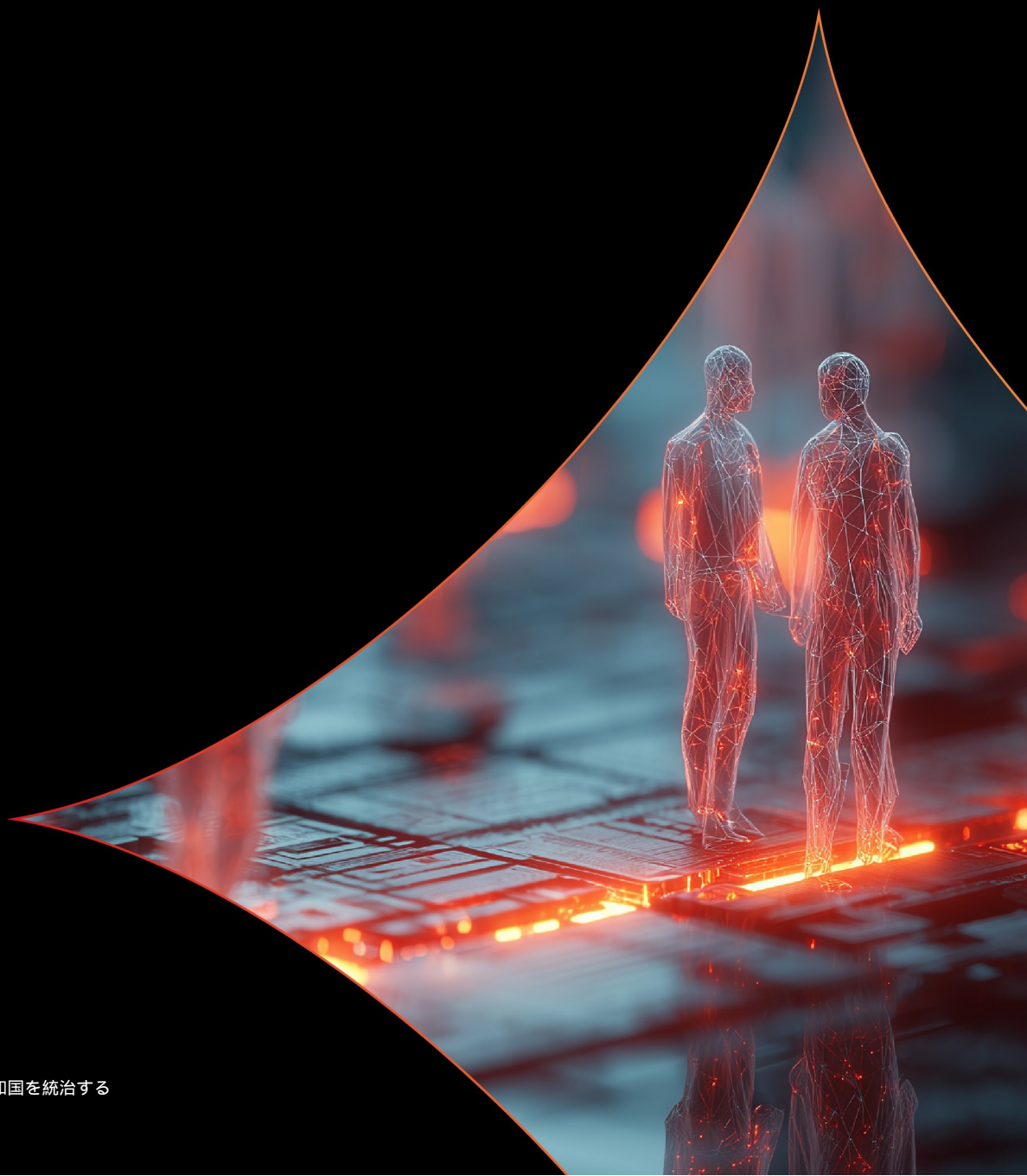
既存のリサーチは所有権とアクセス制御を扱っていますが、エージェントのサプライチェーン検証までは踏み込んでいません。TrendAI™は現在、次のものを開発しています。AIエージェント版の「ソフトウェア部品表」（モデルの出所、学習データの系譜、プロンプトアーキテクチャ、ツールアクセスの文書化）、展開前のスキャンおよび認証プロセス、市民権レジストリ（承認済みのすべてのAIエージェントについて、その権限、所有者、審査状況を一元管理するインベントリ）、そして失効メカニズム（侵害されたAIエージェントを「国外追放」する、すなわち市民権を剥奪する方法）です。

パート3

国境管理：組織に入るものを検証する

エージェントやデータが国境を越えて入ってくるとき、どのような検証や本人確認が行われているのか。

共和国において、国境管理は防衛の第一線です。誰も資格情報を提示せずに入国することはできず、物品は入国前に検査されます。AIエージェントの共和国では、組織の境界を越えるすべてのエージェント、すべてのデータフィード、すべてのAPI接続が検証されなければなりません。



なぜ重要か：境界における攻撃ベクトル

TrendAI™の研究は、組織の国境を越える3つの主要な攻撃ベクトルを記録しています。以下に、それぞれのベクトルと想定される攻撃を挙げます。

1. AIエージェントハイジャック（既存エージェントの乗っ取り）：

- 認証情報窃取（エージェントプラットフォームへのアクセス）
- APIキーの侵害
- セッションハイジャック
- エージェントプロバイダーへのサプライチェーン攻撃

2. AIエージェントポイズニング（データの不正操作）：

- ソースシステム（メール、Slack、ドキュメント）の侵害
- 学習パイプラインまたはRAGインデックスへのインジェクション
- 人事評価の改ざん
- 偽の会議トランスクリプト

3. シャドーAIエージェント作成（無許可のクローン作成）：

- 不正な管理者による、退職者のAIエージェントの作成
- 削除後のバックアップからのエージェント復元
- 公開データのスクレイピングによる外部クロンの構築

この研究はまた、コンシューマー向けAIを国境上の脅威として特定しています。組織の境界をまたぐデータ漏洩を生む3つのベクトルは、次のとおりです。

ベクトル	内容	リスク
シャドーAIによる漏洩	従業員が個人のAIを企業データとともに利用する	データがコンシューマー向けAIのメモリに取り込まれる
個人アカウントの侵害	従業員の個人のGoogle／OpenAIアカウントへの攻撃	コンシューマー向けAI経由で業務コンテキストにアクセスされる
プラットフォーム横断の推論	複数のコンシューマー向けAIのデータを組み合わせる	無許可のクローン作成に向けた完全なプロファイルが構築される

表3. コンシューマー向けAIプラットフォームを介したデータ漏洩のベクトルとリスク

こうしたリスクは、根本的な現実をあらためて突きつけます。キルチェーンは、システムの奥深くから始まるものではありません。何かが国境を越えたその瞬間から始まり得るのです。その境界において強固なアイデンティティ管理と検証の仕組みを確立できるかどうか、偵察や初期アクセスを早期に阻止できるか、それとも全面的な侵害へと進行させてしまうかの分かれ目になります。

キルチェーンは国境から始まる

TrendAI™リサーチの「AIエージェント侵害キルチェーン」は、偵察と初期アクセスから始まります。いずれも、国境を越えるイベントです。

- **フェーズ1（偵察）**：標的組織の特定、価値の高い従業員のマッピング、AIエージェントプラットフォームの列挙、個人情報（PII）の収集
- **フェーズ2（初期アクセス）**：エージェントプラットフォームの認証情報を狙うフィッシング、APIの脆弱性の悪用、管理者アカウントの侵害、連携先（Slack、メール、カレンダー）へのサプライチェーン攻撃

国境管理がフェーズ1またはフェーズ2の段階で侵入を捕捉できれば、その先のフェーズは発生しません。

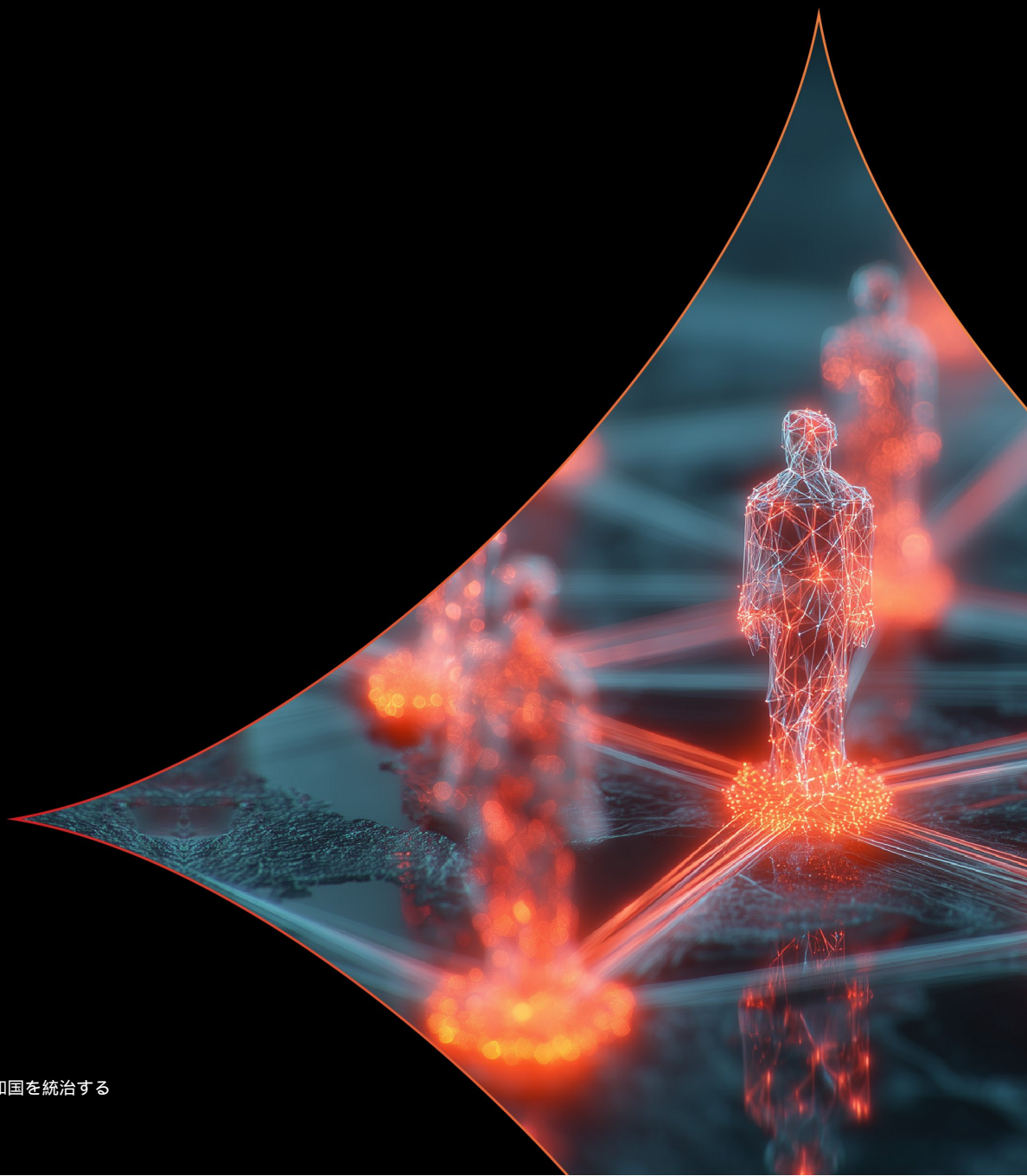
最大のギャップは、AIエージェントのためのゼロトラストセキュリティアーキテクチャです。既存のリサーチは、行動ベースライン、ライブネス相関、異常検知といった展開後の検知手段を提供していますが、流入前の検証ゲートには対応していません。TrendAI™は、(a) 境界におけるAIエージェントのアイデンティティ検証、(b) データ検証ゲート、(c) 継続的検証（ゼロトラストの原則は「決して信頼せず、常に検証する」）、(d) AIエージェント向けマイクロセグメンテーションを開発しています。これにより各エージェントの影響範囲が限定され、1体のエージェントの侵害が連鎖的に拡大する事態を防ぎます。

パート4

抑制と均衡：エージェントのアクションにおける職務分掌

資金の移動や重要な意思決定においては、依頼と実行を別々の主体が担わなければならない。支払いを依頼する人と小切手を切る人が別人でなければならないという、会計の原則と同じである。

共和国では、抑制と均衡の仕組みによって、政府のどの機関も歯止めのない権力を蓄積できないようになっています。AIエージェントの共和国においても、重要なアクション、とりわけ金銭に関わるアクションについて、単一のエージェントがエンドツーエンドの権限を持つべきではありません。



信頼レイヤーの悪用

信頼関係に関する知識を持つAIエージェントが侵害されると、その情報は効率的かつ効果的に悪用され得ます。たとえば、どの同僚なら確認せずに応じるかを見極め、誰の「緊急依頼」なら即座に対応されるかを把握できます。そのうえで、統制を迂回する非公式な承認経路を悪用し、関係の履歴をソーシャルエンジニアリングに利用することもできてしまいます。

TrendAI™リサーチの「クイックスタート検知」には「権限昇格」ルールが含まれており、AIエージェントが従業員の通常の権限レベルを超えた承認や認可に使われた場合を追跡します。しかし、検知はあくまで事後対応です。抑制と均衡は、構造として組み込まれなければなりません。

TrendAI™は、AIエージェントのアクションに対する職務分掌フレームワークを開発しています。既存のリサーチはリスクを明確に特定しているものの、構造的な解決策までは提示していません。TrendAI™が開発しているのは、次のとおりです。

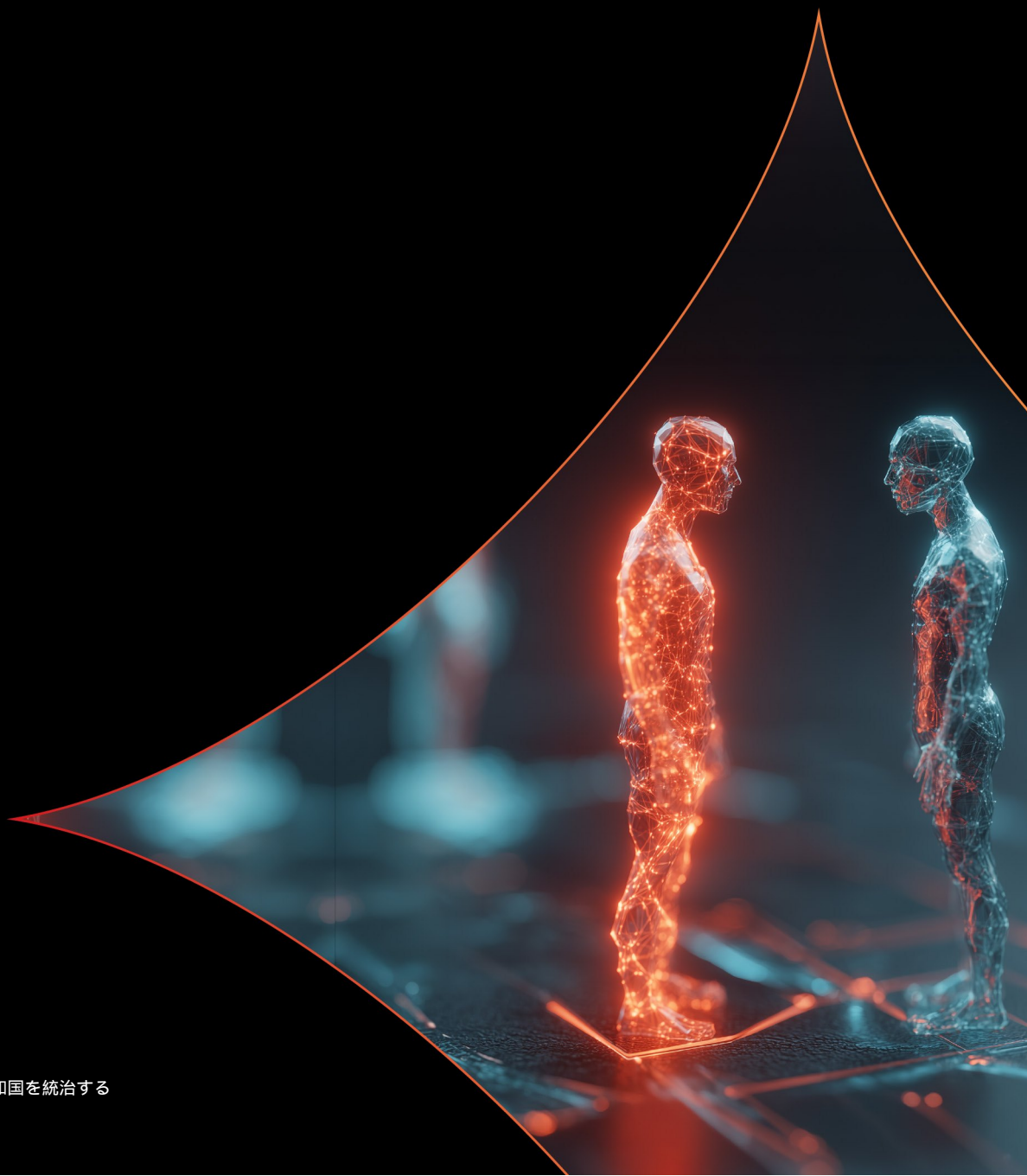
- 複数主体の承認を必要とする重要アクションの分類
- 依頼側のAIエージェントと実行側のAIエージェントを、異なる信頼チェーンを持つ別個の主体としなければならないというアーキテクチャパターン
- どの金額、どのデータ機密性レベル、どの権限レベルから、第2のAIエージェントまたは人間の関与を必須とするかを定めるエスカレーションしきい値
- 支払いの依頼と実際の支出に別々の承認済み署名者2名を求める会計原則になぞらえた、「連署（counter-signing）」のプロセス

さらに、ヒューマンインザループの要件も欠かせません。AIエージェントの能力にかかわらず、常に人間の承認を必須とする意思決定をあらかじめ定義しておくことです。

マルチエージェント合意：重要な意思決定のために

重大な意思決定には、3体を超えるエージェントの合意が必要である。重要なリサーチでは、同じ問いを2つを超えるLLMモデルで処理し、クロス検証しなければならない。

共和国では、大型の調達には競争入札が求められます。競争なしに契約を勝ち取るベンダーは存在しません。重要な法案の成立には、複数の議院の合意が必要です。この競争による検証の原則は、AIエージェントの意思決定にも拡張されなければなりません。



なぜ重要か

TrendAI™の研究は、たった1体の侵害されたAIエージェントが壊滅的な被害をもたらし得ることを実証しています。ハイジャックされた、あるいはポイズニングされたAIエージェントの意思決定は、信頼された発信元から届くため、一見正当に見えます。組織が重要な意思決定を単一のAIエージェントの出力に依存しているなら、それは単一障害点を抱えているということです。

わずかに誤った出力を生み出し得る攻撃ベクトルは、3つあります。

- **ハイジャック**：攻撃者がAIエージェントを掌握し、意図的に偏った出力や悪意のある出力を生成する
- **ポイズニング**：AIエージェントの学習データが改ざんされ、「ほぼ正常だが、バイアスやバックドアが埋め込まれた」出力が生成される
- **思考様式の不正操作**：AIエージェントの意思決定モデルへの問い合わせを通じて、ビジネス上の意思決定が抽出または誘導される

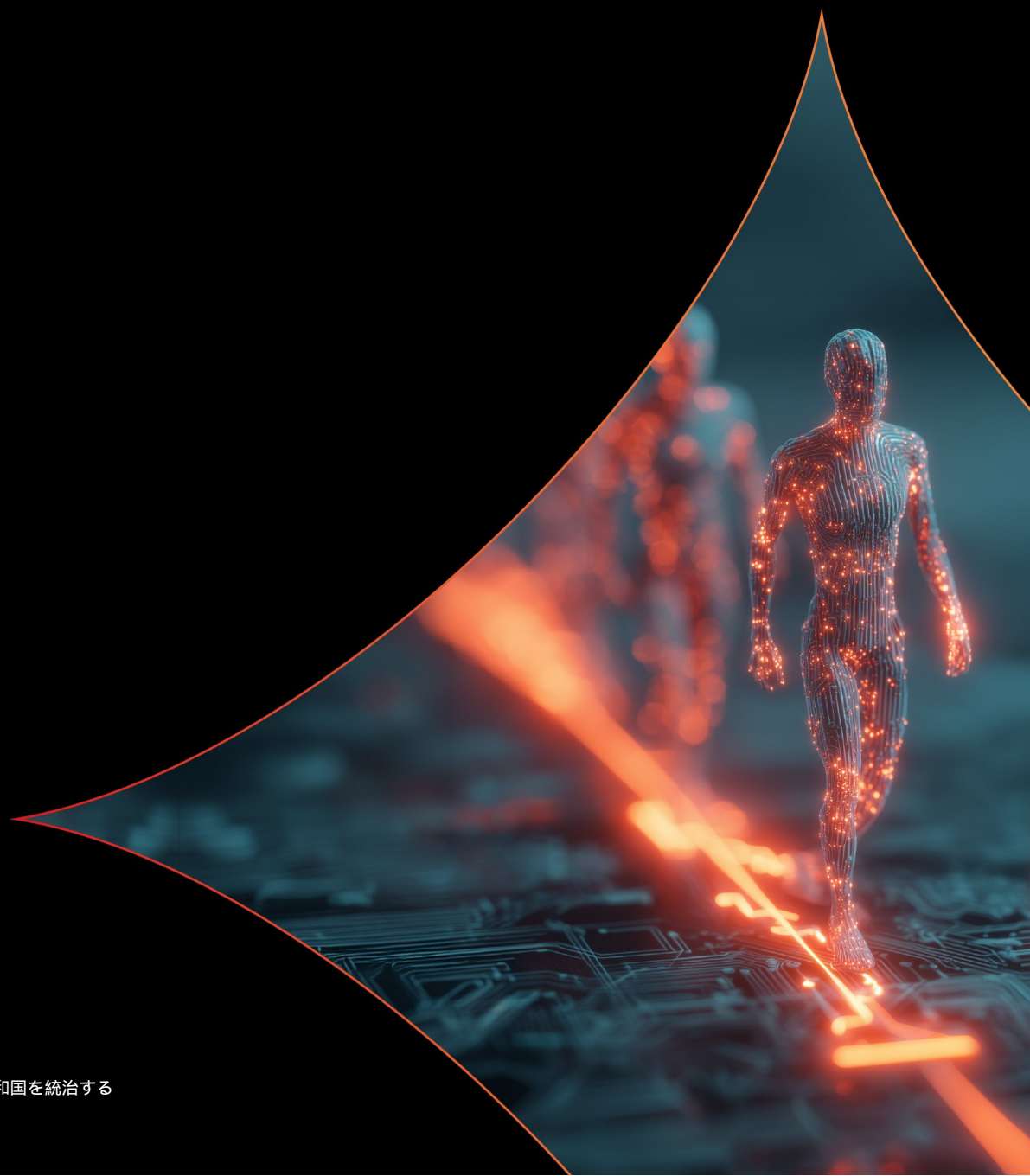
3つのいずれの場合も、出力は正当に見えます。もっともらしく見えて実は侵害されている推奨を単一のエージェントが提示すれば、独立した裏付けを求める構造的な仕組みがない限り、疑われることなく受け入れられてしまうでしょう。

TrendAI™は、マルチエージェント合意フレームワークを開発しています。そこには、マルチエージェント合意を必要とする意思決定の分類、アーキテクチャパターン、モデル間のクロス検証、そしてパフォーマンスとコストのトレードオフ（マルチエージェント合意は高コストで低速です）が含まれます。

監査証跡：全エージェント通信のログ 記録の義務化

すべてのエージェントは、あらゆる通信とアクションの明確な記録を保持し、完全な監査証跡を残さなければならない。

共和国において、透明性と説明責任は譲れないものです。政府のあらゆる行為は記録され、あらゆる金銭取引には証跡が残り、あらゆる公式な通信はアーカイブされます。AIエージェントの共和国にも、同じことが求められます。



なぜ重要か：リサーチが示す明確な根拠

これは、最も強力な裏付けを持つガバナンスの柱です。TrendAI™のリサーチはその必要性を明示的かつ強力に論じており、TrendAI™の検知フレームワークは、包括的なログが存在するという前提の上に成り立っています。監査証拠がなければ、次の検知カテゴリーはいずれも機能しません。

カテゴリー	主要なシグナル	ログが必要な理由
行動ベースライン	クエリパターンの変化、語彙の変化、応答時間の異常	現在の行動を過去のベースラインと比較する必要がある
ライブネス相関	従業員が明らかに別の場所にいる間のAIエージェントの活動	AIエージェントのアクセスログと物理的な入退室記録を突合する必要がある
退職後の活動	退職した従業員のAIエージェントへのクエリ	退職後のあらゆる活動を検知する必要がある
なりすまし	チャンネル間の不整合、権限昇格	すべてのコミュニケーションチャンネルにわたって行動を追跡する必要がある
データ持ち出し	一括エクスポート、埋め込みストアへのアクセス、トークン窃取	データ移動の量とパターンを監視する必要がある

表4. AIエージェントの重要な検知カテゴリーと必要なログシグナル

リサーチからの推奨事項

TrendAI™の推奨事項は明快です。

- 初日からログ記録と監視を実装する
- どのようなAIエージェントが存在し、誰がアクセスできるかを棚卸しする
- インシデント対応プレイブックを作成する

規制の動向も、この方向性を後押ししています。EUの一般データ保護規則（GDPR）、カリフォルニア州プライバシー権法（CPRA）によって改正されたカリフォルニア州消費者プライバシー法（CCPA）、イスラエルのプライバシー保護法改正第13号、そして日本のAI推進法（正式名称「人工知能関連技術の研究開発及び活用の推進に関する法律」）のいずれにも、AIシステムの活動記録を提出できることを前提としたコンプライアンス要件が定められています。

AIエージェントの共和国では、エージェント同士が相互に通信する場面が増えていきます。TrendAI™は、人間とエージェントのやり取りに焦点を当てた既存のリサーチを発展させ、監査証拠の範囲をエージェント間通信へと拡張します。ログ記録の要件も、すべてのエージェント間通信チェーンをカバーするように拡張されます。

さらに、エージェント間の監査証拠に何が含まれるべきかも定義します。送信者、受信者、内容、タイムスタンプ、決定の結果、行使された権限です。ログの保護と完全性に関する要件も策定します。ログは、改ざんできない状態にあってはじめて意味を持つからです。

結論

組織は今、知識、人格、権限、信頼を備えて活動するワークフォースを構築しつつあります。この新しいワークフォースには、新しいルールが必要です。決定論的なソフトウェアから自律的なAIエージェントへの移行は、単なる技術の転換ではなく、ガバナンスの転換なのです。

本ホワイトペーパーで提案した6つのガバナンスの柱、すなわち憲法、市民権、国境管理、抑制と均衡、マルチエージェント合意、監査証跡は、この新しい共和国を管理するためのフレームワークとなります。これらの柱の中には、既存の研究によって十分に裏付けられているものもあれば、業界がこれから開発していくべき最先端のガバナンス概念もあります。

ただし、統治するという行為そのものに、選択の余地はありません。TrendAI™の研究が結論付けているとおり、2027年末までに、AIエージェントの侵害は認証情報の窃取よりも深刻な問題となるでしょう。盗まれたAIエージェントは、リセットできないからです。ガバナンスフレームワークを構築すべきなのは、共和国が住民で満たされてからではなく、その前です。

参考文献：

Kirill Gelfand, Vladimir Kropotov, Fyodor Yarochkin, and Robert McArdle. (2026年3月26日). TrendAI™. 「[Unconventional Attack Surfaces Identity Replication via Employee Digital Twins](#)」 (英語) . 2026年4月13日閲覧.

Vladimir Kropotov, Fyodor Yarochkin, and Kirill Gelfand. (2026年2月10日). TrendAI™. 「[AI Skills as an Emerging Attack Surface in Critical Sectors: Enhanced Capabilities, New Risks](#)」 (英語) . 2026年4月13日閲覧.

TREND MICRO

本書に関する著作権は、トレンドマイクロ株式会社へ独占的に帰属します。

トレンドマイクロ株式会社が書面により事前に承諾している場合を除き、形態および手段を問わず本書またはその一部を複製することは禁じられています。本書の作成にあたっては細心の注意を払っていますが、本書の記述に誤りや欠落があってもトレンドマイクロ株式会社はいかなる責任も負わないものとします。本書およびその記述内容は予告なしに変更される場合があります。

本書に記載されている各社の社名、製品名、およびサービス名は、各社の商標または登録商標です。

〒160-0022

東京都新宿区新宿 4-1-6 JR新宿ミライナタワー

<https://www.trendmicro.com>

トレンドマイクロはサイバーセキュリティのグローバルリーダとしてデジタル情報を安全に交換できる世界の実現に貢献します。私たちの革新的なソリューションはデータセンター、クラウド、ネットワーク、エンドポイントにおける多層的なセキュリティをお客様に提供します。

当社のリーダシップの根幹であるトレンドマイクロリサーチは、多くのエキスパートに支えられています。それは最新の脅威を発見し、重要なインサイトを公に共有し、サイバー犯罪の防止を支援することに情熱を注ぐ人材です。当社のグローバルチームは、日に数百万もの脅威を特定し、脆弱性の開示を先導し、標的型攻撃・AI・IoT・サイバー犯罪等における革新的な研究結果を公表しています。私たちは次に来る脅威を予測し、セキュリティ業界が進むべき方向を示しうる示唆に富んだ研究成果を提供するため、継続的に取り組んでまいります。

© 2026 Trend Micro Incorporated. All Rights Reserved.



さらに詳しい情報をお求めの方へ

TrendAI | Resources and Insights



TrendAI™は、AIセキュリティのグローバルリーダーであり、Trend Microのエンタープライズ事業部門として、AIの完全な可視化と統合されたセキュリティを提供します。それにより、組織に確かな信頼をもたらし、イノベーションを促進し、リスクを排除します。

TrendAI™は、185の国と地域にわたる大手企業や政府機関から信頼を得ており、アイデンティティからインフラストラクチャ、データに至るまで、組織全体を保護します。

AI Fearlessly.

詳細はこちら：trendaisecurity.com

Copyright © 2026. Trend Micro Incorporated. All rights reserved. [REP00_Governing_the_Republic_of_AI_Agents_100626US]