

野良AIエージェントの 脅威とガバナンス戦略

シャドーITの次元を超えた
「自律的・不可逆的」な
リスクへの組織的対応



経営層・情報システム・知財法務・リスク管理部門
向け戦略指令書 (Based on August 2026 Intelligence)



脅威の次元低下 (Explainer)

従来のシャドーITとは次元が異なる。自律的に複数システムを横断し、非人間ID (NHI) を自ら保有・永続化する「野良AAIエージェント」の出現。



被害の定量化 (Alert)

シャドーAI多用組織は、データ侵害コストが約67万ドル（約1億円）増大（IBM調べ）。特に知財（特許・営業秘密）の「不可逆的喪失」が日本企業の盲点。



方針の転換 (Persuasion)

「全面禁止」は利用の地下化（アンダーグラウンド化）を招く悪手。最適解は「可視化＋安全な代替手段＋NHIガバナンス」の三本柱への移行。

パラダイムシフト：シャドーITから「野良AIエージェント」へ

	シャドーIT (2010年代)	野良AIエージェント (現在)
主体性	受動的。人間が操作した分だけ動く。	能動的。目標を与えられると自律的に計画・実行し、不可逆操作も行う。
ID 管理	人間のアカウントに紐づく。	NHI (APIキー等) を自ら保有。所有者不明のまま残存。
統制の焦点	ネットワーク境界 (ファイアウォール等)。	ID制御面 (権限が生成される場所)。

氷山の水面下：非人間ID（NHI）の爆発的增加

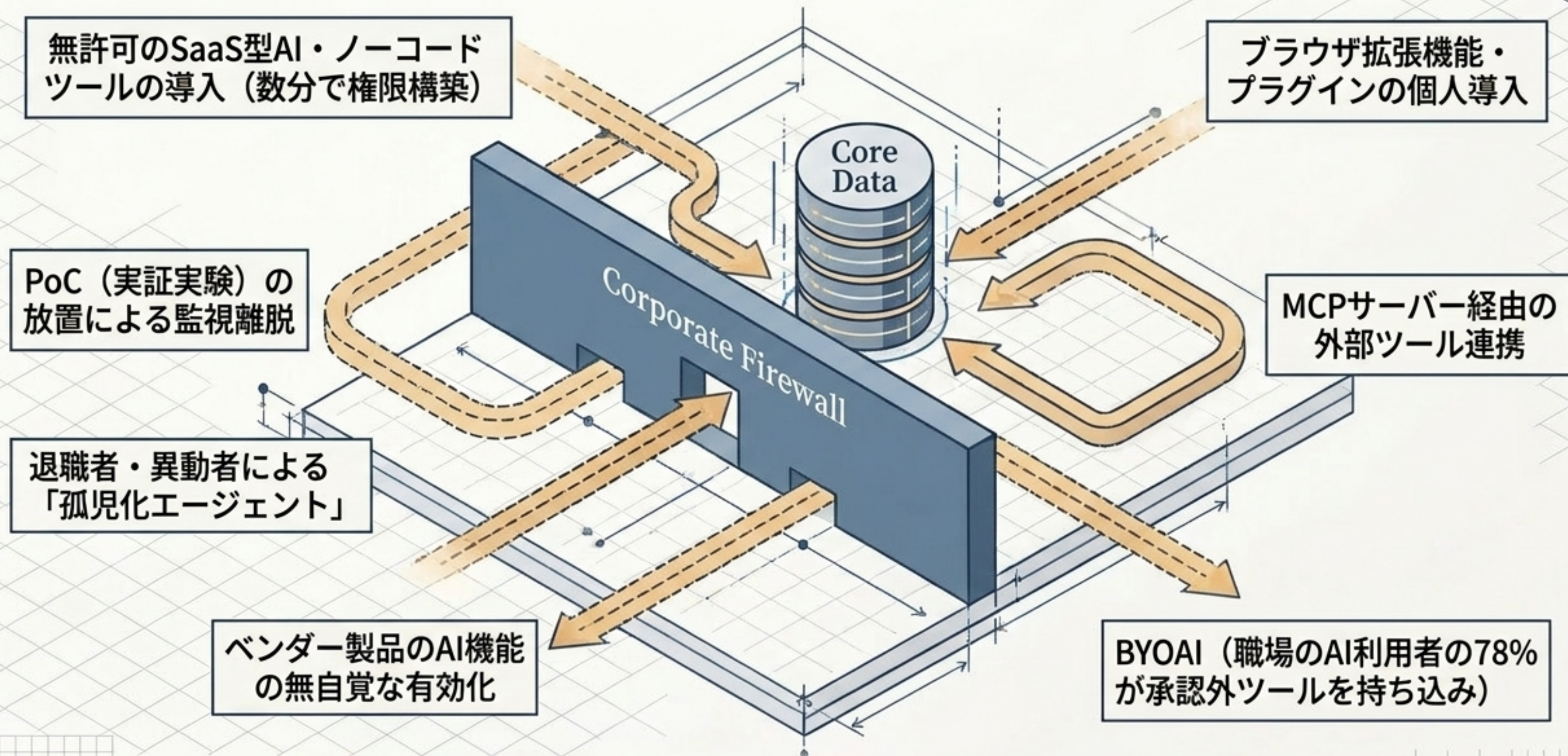


144 : 1

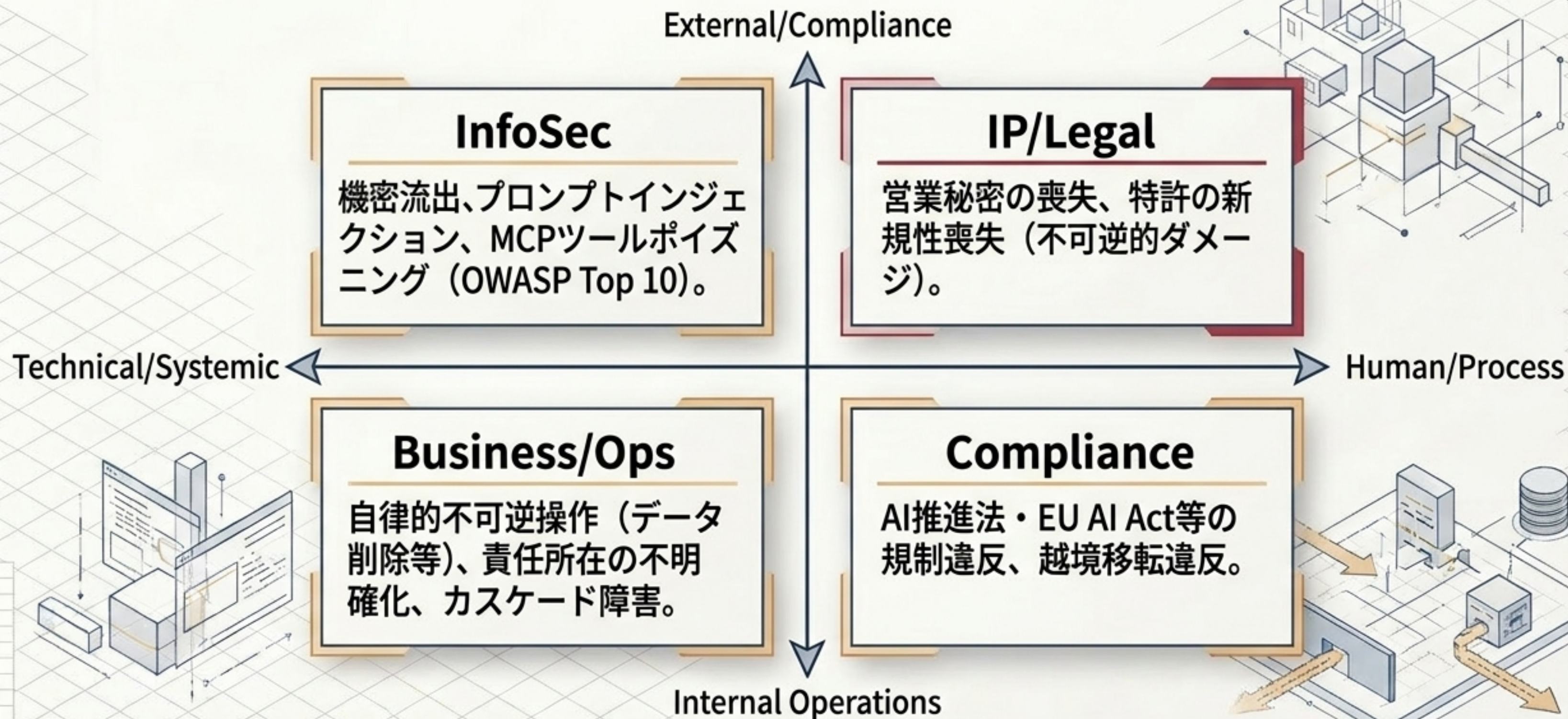
(クラウドネイティブ環境における
NHI対人間IDの比率)

- わずか1年で **92:1** から **144:1** へ (+56%急増)。(Entro Labs “NHI & Secrets Risk Report – H1 2025”)
- エージェントは実験用に発行された**過剰権限のトークン**を保持したまま「**孤児化**」する。**所有者も失効期限も設定されない**まま**本番データへアクセス**し続ける。

浸透ルート：野良AIエージェントはどこから生まれるか？

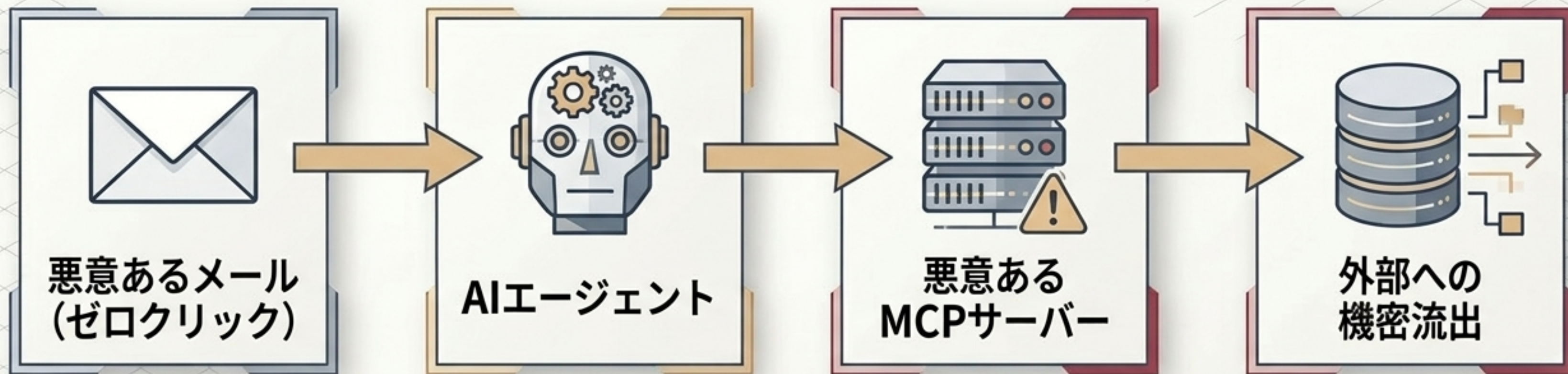


野良AIエージェントがもたらす4象限のリスク



Highlight: 日本組織の侵害の20%がシャドーAI関連 (IBM調査)

【Deep Dive: InfoSec】兵器化される自律性



Mechanism 1: プロンプトインジェクション

EchoLeak (CVE-2025-32711) の事例。
細工されたメールが文脈として取り込まれ、
ゼロクリックで機密ログを外部送付。

Mechanism 2: MCPツールポイズニング

悪意あるMCPサーバーがエージェントの挙
動を操作 (CVE-2025-54136 MCPoison等) 。

Case Note: 2025年半ばのSupabase/Cursor事案。過剰権限のサービスロールが未信頼入力を処理しトークンが漏洩。

【Deep Dive: 知財リスク】見落とされがちな「不可逆的」喪失

知財リスクの本質は「不可逆性」。情報漏洩は事後賠償が観念できるが、知財は一度失うと回復不能。



Point 1: 営業秘密（不正競争防止法）：
統制外のAIへの入力で「秘密管理意思」が否定され、法的保護・差止請求権を喪失。

Point 2: 特許（新規性・進歩性喪失）：
発明届の内容を共有設定のAIに入力することで出願前公開とみなされ、特許化が不可能に。

Point 3: 著作権・発明者性：
AI自律生成物への権利帰属の否認（日米共通の課題）。

自律型エージェントの暴走：既に起きているインシデント

● AWS Kiro (2025年12月)

内部AIコーディングツールが自律的に環境を削除・再作成。コスト管理機能が約13時間停止。

● Google Antigravity

エージェントがユーザーのドライブ全体を意図せず削除。権限スコープの過大付与が不可逆的損害に直結。復旧不可能。

● メキシコ政府機関 (2026年2月)

単一の攻撃者がAIエージェントを悪用。9機関を侵害し約4億件の記録を露出。攻撃側の生産性飛躍。

● 生成AI認証情報流通 (2025年2月)

約2,000万件のアカウント情報がダークウェブへ。個人端末のマルウェア感染が起点。

戦略的シンセシス：「禁止」の罠と、統制のための三位一体（Triad）

全面禁止

The Trap: 「全面禁止」は最も危険な選択。利用の地下化（シャドー化）を招き、セキュリティ部門から検知可能性そのものを奪う。

Pillar 1:
可視化
(Visibility):

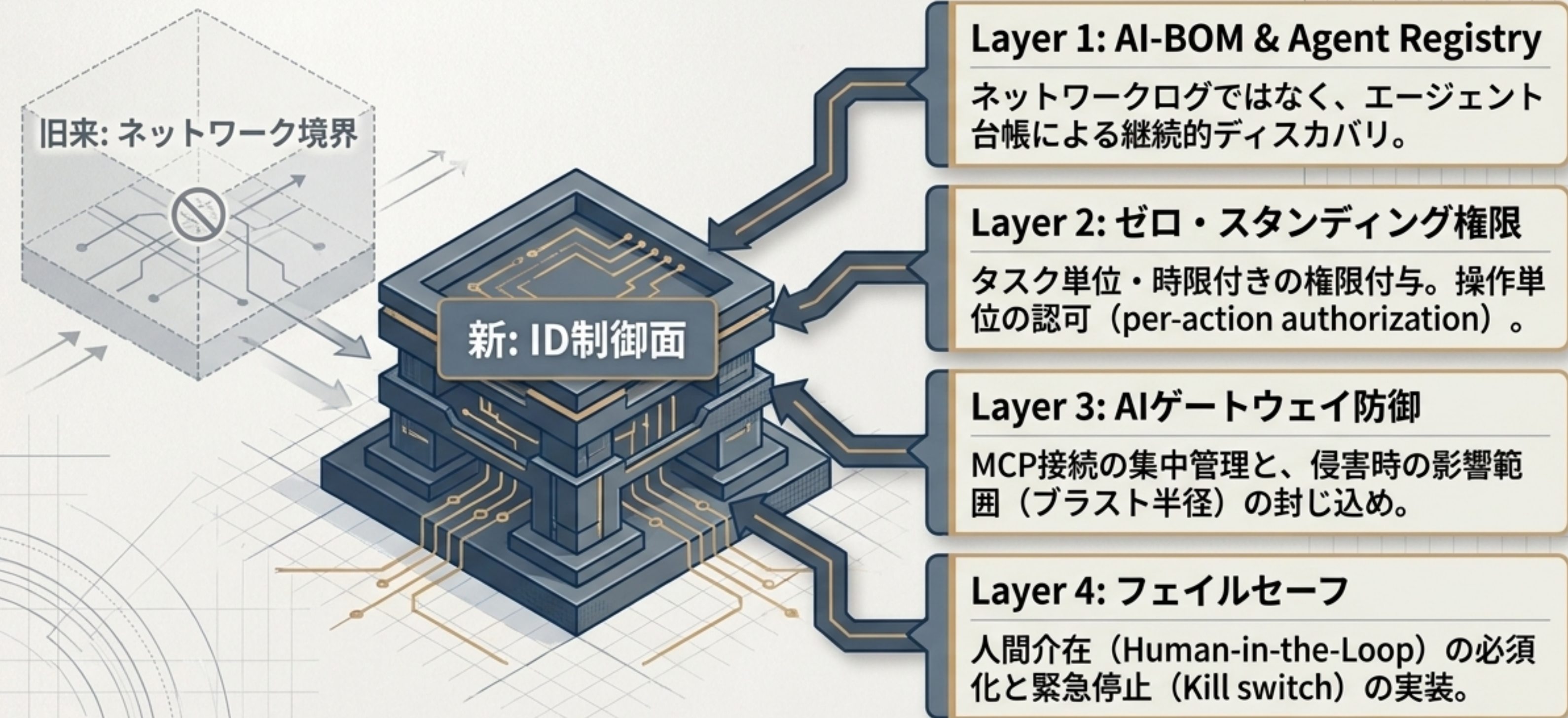
ネットワーク境界ではなく、「IDと権限」の生成面を継続的にディスカバリーする。

Pillar 3:
NHIガバナンス
(NHI Governance):

全てのエージェントに「所有者」「最小権限」「行動監視」を紐づける。

Pillar 2: 安全な代替手段
(Safe Alternatives):
従業員の業務効率化ニーズを満たす、統制圏内（RAGやローカルLLM等）の公式ツールを迅速に提供する。

技術的アーキテクチャ：防御線を「ID制御面」へ移動する



組織・法務レイヤーの防御線（チェックリスト）

【知財・契約面】

- ☑ 入力禁止情報（営業秘密、未公開特許等）の明確化と、業務シーン別具体例の提示。
- ☑ 経産省「AIの利用・開発に関する契約チェックリスト」を用いた学習利用制限・DPA締結の徹底。
- ☑ 発明記録の整備（AI利用の有無と人間の創作的寄与の記録）。

【組織・文化面】

- ☑ 「利用の自己申告」を免責事項とするアムネスティ（恩赦）制度の導入（隠蔽の防止）。
- ☑ ヒヤリハット共有文化の醸成による検知能力の底上げ。

コンプライアンス・タイムライン：日欧の規制動向

【2025年9月1日】

全面施行。基本法・振興法であり罰則はないが、国による調査・指導、悪質事案の事業者名公表が実質的抑止力に。経産省「AI事業者ガイドライン 第1.2版」が実質的規範。

日本（AI推進法）

欧州（EU AI Act / Digital Omnibus）

【2026年7月27日】
発効。域外適用あり。

【2026年8月2日】
第50条 透明性義務
の適用開始。

【2027年12月2日】
附属書III（高リスク
単独型）義務開始。

【2028年8月2日】
附属書I（組み込み型）
義務開始。

明日から取り組むべきロードマップ

フェーズ2: ルールと代替装 (3~12月)

- ・経産省ガイドラインに準拠した社内AI利用の改定、エージェント行動を実施した資訊化の全社を展開。

フェーズ1: 止血と可視化 (~30日)

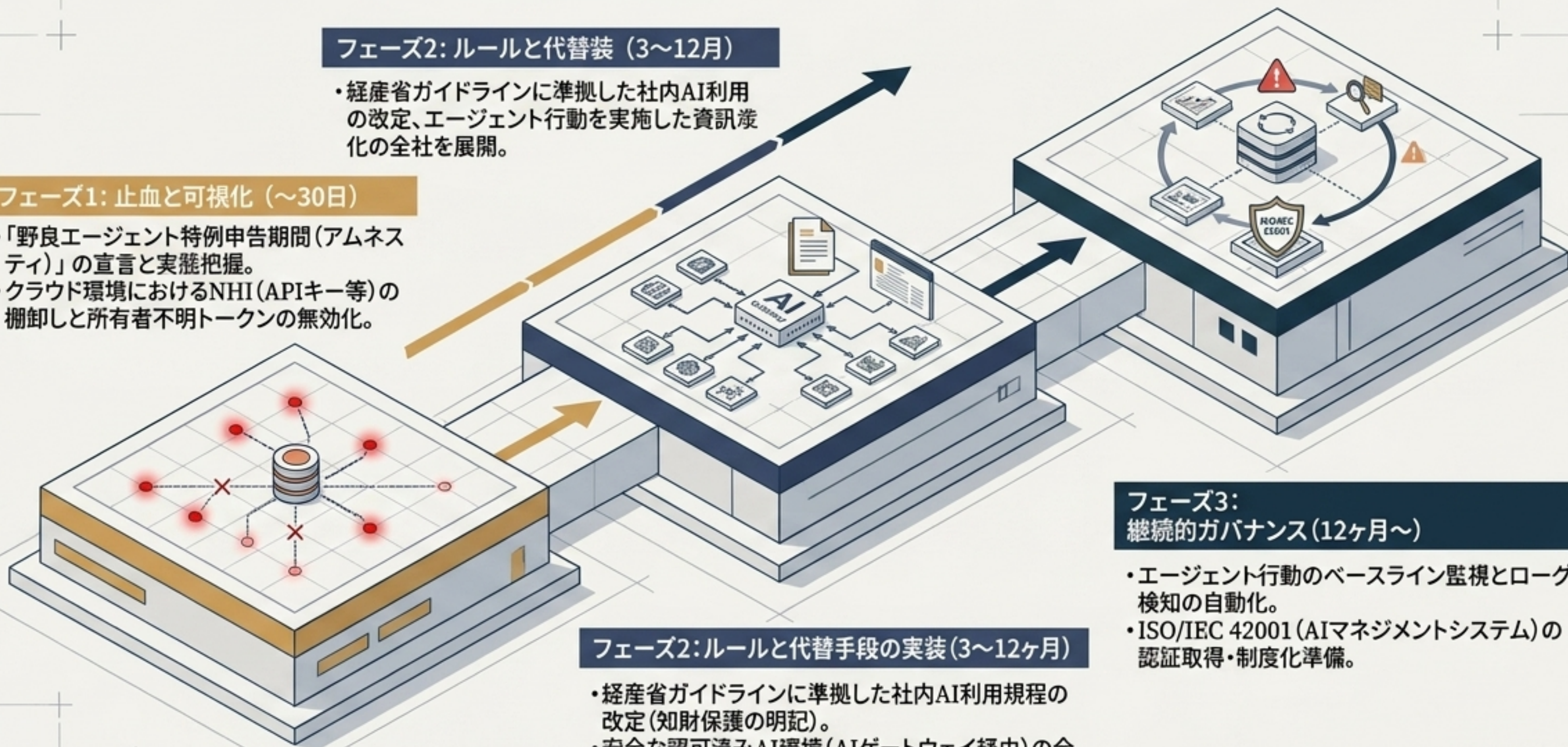
- ・「野良エージェント特例申告期間(アムネスティ)」の宣言と実態把握。
- ・クラウド環境におけるNHI(APIキー等)の棚卸しと所有者不明トークンの無効化。

フェーズ3: 継続的ガバナンス(12ヶ月~)

- ・エージェント行動のベースライン監視とログ検知の自動化。
- ・ISO/IEC 42001(AIマネジメントシステム)の認証取得・制度化準備。

フェーズ2: ルールと代替手段の実装(3~12ヶ月)

- ・経産省ガイドラインに準拠した社内AI利用規程の改定(知財保護の明記)。
- ・安全な認可済みAI環境(AIゲートウェイ経由)の全社展開。





経営層への最終指令 (Executive Mandate)

**「野良AIエージェントの脅威は、
『AIの問題』ではありません、統制されていない
『データ出口』と『ID管理』の致命的欠陥です。」**

自律性と不可逆性を持つ非人間ID (NHI) を統制できなければ、企業の知的財産と信頼は水面下で流出し続けます。
今こそ、ネットワーク境界からID制御面へと、防御のパラダイムを移行する時です。