

# 麒麟「AI研究員」の実像と死角

生成AI時代のR&D知財戦略と次の一手



# Executive Summary

## ① 虚像と実像 の分離

報道の「AI研究員」は自律的  
発明主体ではない。公式発表  
は研究者の創造性を引き出す  
「伴走」として位置づけている。

## ② 新規性の所在

「業務効率化」から「知的  
創造活動そのもの」への適用  
対象の移動。技術スタックや  
成果目標は非開示。

## ③ 最大のリスク

上流（仮説生成）の加速に対  
し下流（他社権利）との接続  
が未確認。知財部門がボトル  
ネック化する構造的死角。



# Expectation vs. Reality : 虚像と実像の分離

メディアの熱狂

# AI研究員

AIが自律的に研究・発明を行う主体であるかのような誤解。

公式発表の実像 (2026年8月6日)

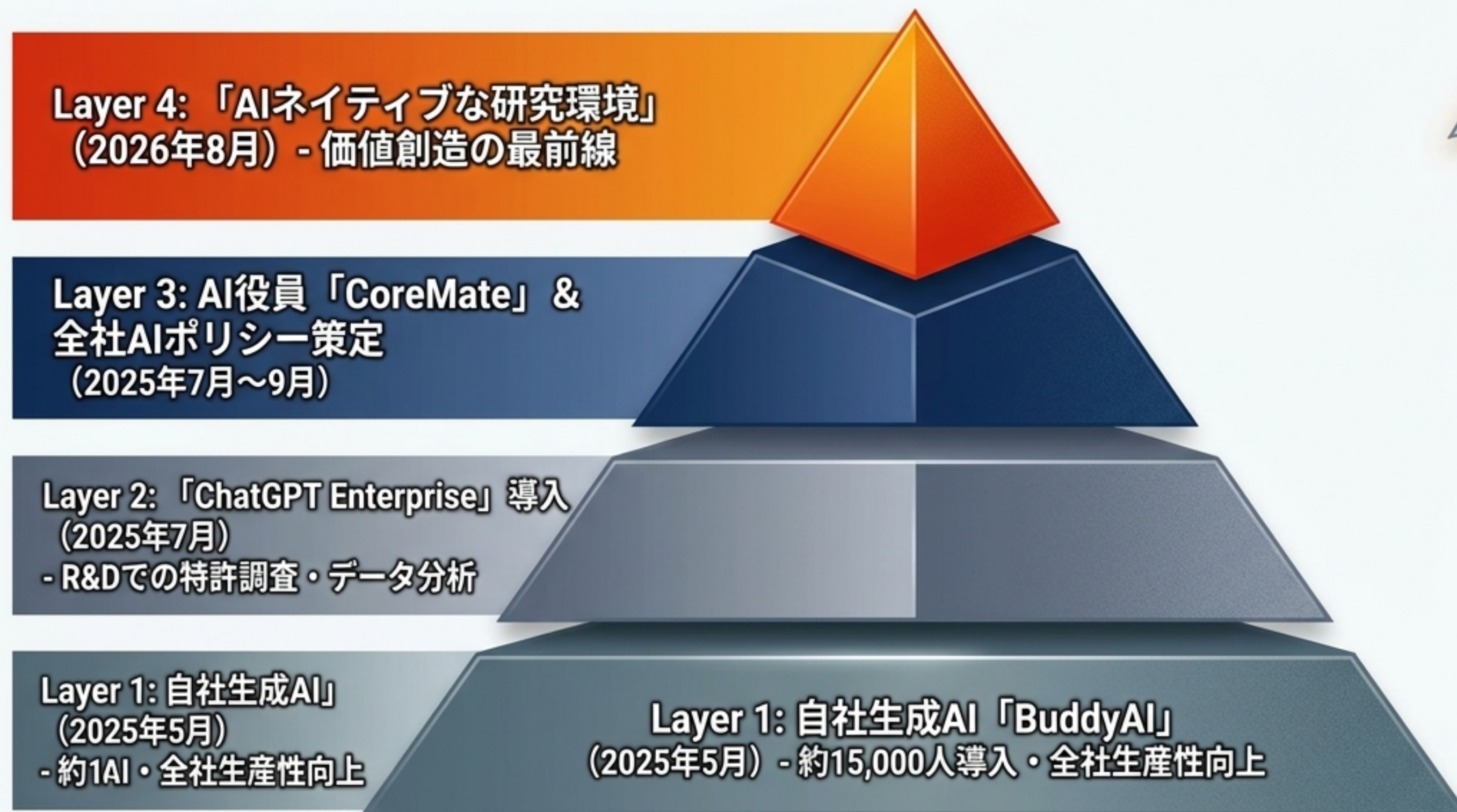
代替ではなく、研究者を  
主役とする「伴走者」

GenerativeX社との協業による、  
研究活動への常時組み込み

対象業務は「仮説形成、情報探索、  
壁打ち」。物理的実験の自動化への  
言及はなし

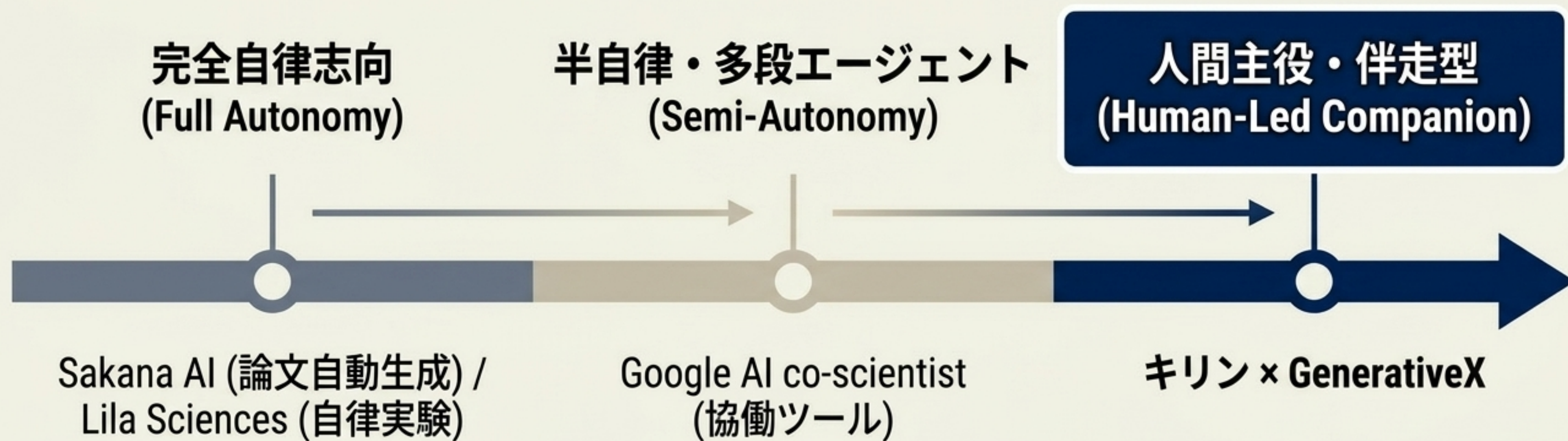


# Strategic Blueprint : キリンのAI・DX戦略スタック





# The Autonomy Spectrum : 自律度におけるポジショニング



**Key Insight:** 麒麟は技術的な「完全自律」の先進性よりも、既存の企業R&D現場への「定着と摩擦ゼロ」を意図的に選択している。



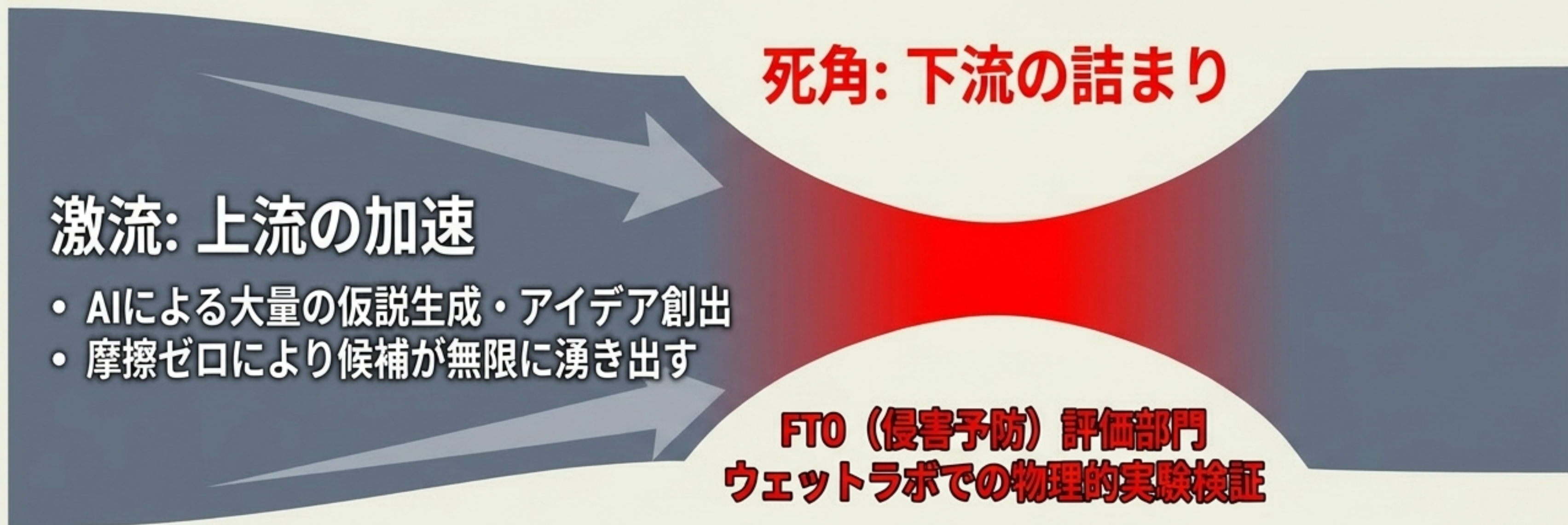
# Industry Comparison : 国内業界AIアプローチの比較

| 食品・飲料業界                                                                                                              | 化学・素材業界 (MI)                                                                                               | キリン (本件)                                                                                                                  |
|----------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------|
| <ul style="list-style-type: none"><li>主目的: マーケティング、業務効率化</li><li>実例: アサヒ (SNS画像からの新商品案生成)、サントリー (独自AIガウディ)</li></ul> | <ul style="list-style-type: none"><li>主目的: マテリアルズ・インフォマティクス、物性予測</li><li>実例: 少実験・高精度化による材料探索の自動化</li></ul> | <ul style="list-style-type: none"><li>主目的: 研究プロセス (思考) の支援、仮説形成</li><li>特徴: データ予測ではなく「研究者の知的創造の拡張」に軸足を置く独自ポジション</li></ul> |

食品発酵領域にMIの潮流を持ち込みつつ、あくまで「人間の思考支援」に特化。



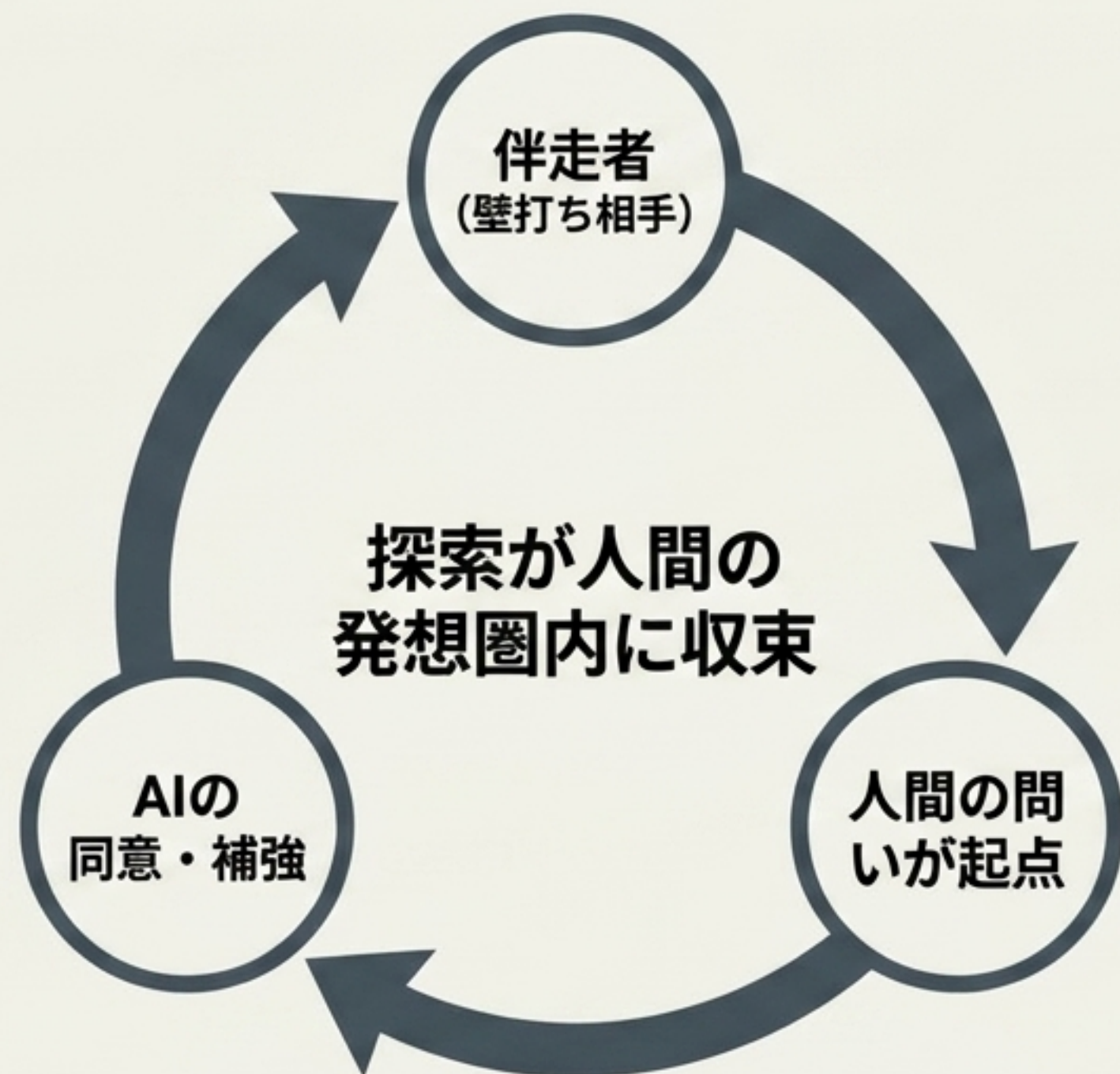
# The Deep Blind Spots : 構造的ボトルネック



**最大の構造的リスク: 上流の仮説生成だけを加速させても、下流の知財調査と実験室の処理能力が変わらなければ、R&Dプロセス全体がスタックし知財部門が機能不全に陥る。**



# The Paradox of Companion AI : 思考支援のパラドックス



確証バイアスの増幅装置 (レッドチームの不在)

思考の分断  
(Friction)



効率化  
(Efficiency)

## 摩擦と偶発性の喪失

「摩擦の除去＝効率化」は、同時に異分野との偶発的衝突（セレンディピティ）という創造性の源泉を奪うリスクを孕む。



# Legal Reality：法的・知財的現実の収斂

米国: 2025年11月28日

USPTO改訂ガイダンス：AIは道具。  
伝統的な着想（Conception）基準  
を適用。

日本: 2026年3月4日

最高裁決定（ダバス事件）：現行特許法上  
の発明者は自然人に限られると確定。

日米ともに  
「AIは道具であり、  
発明者は自然人」で  
完全収斂。

🔍 最大の構造的リスク

実務の焦点は「AIが作ったか」ではなく、  
「人間の着想を事後にどう証跡として証  
明するか」へ移行。



# Governance Dilemma：情報共有のジレンマ

## AI探索履歴の組織共有

### ルートA: 営業秘密

- ・ 厳格な秘密管理性が必要
- ・ 矛盾：共有範囲の拡大と管理の両立は困難

### ルートB: 限定提供データ

- ・ 共有を前提とするデータ管理
- ・ 提供条件と契約の厳密な整備が必要

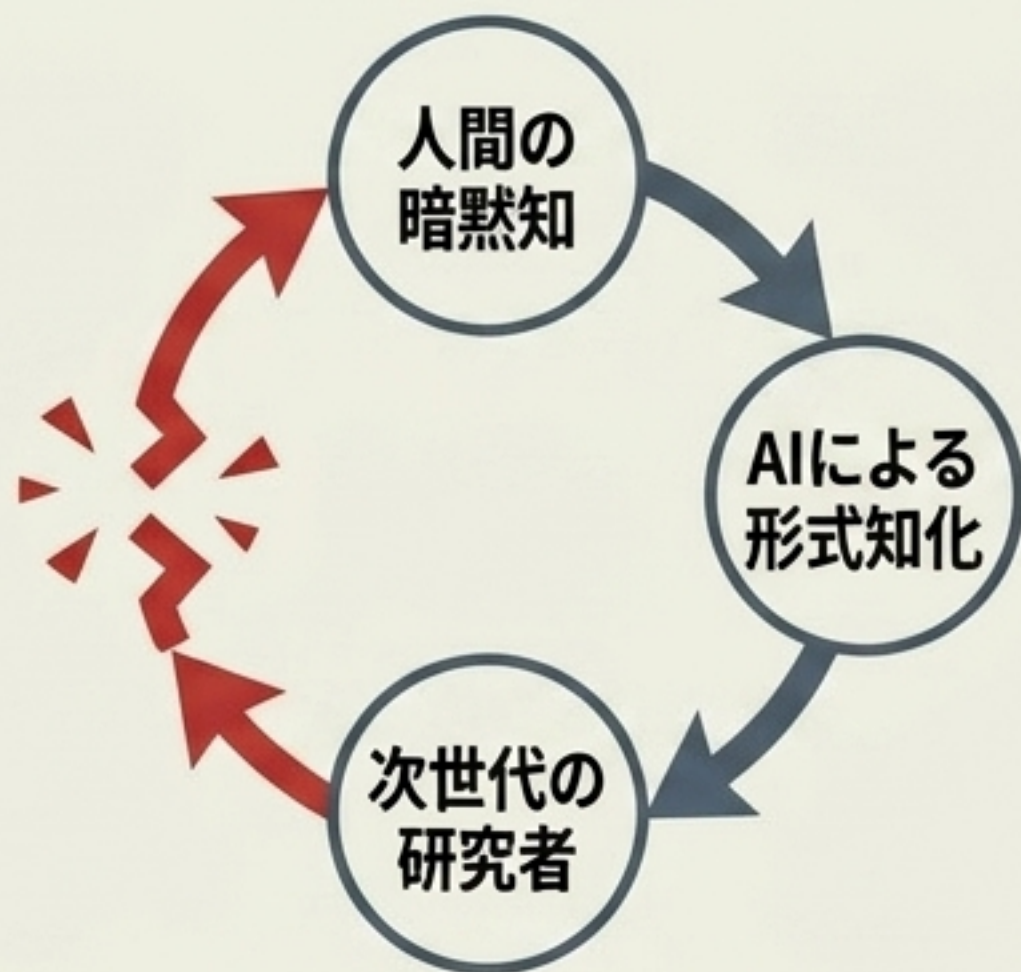


外部ベンダー経由のRAG環境において、全社AIポリシーだけで保護要件を満たせるか。情報の性質に応じた「区別とアクセス制御」の再設計が急務。



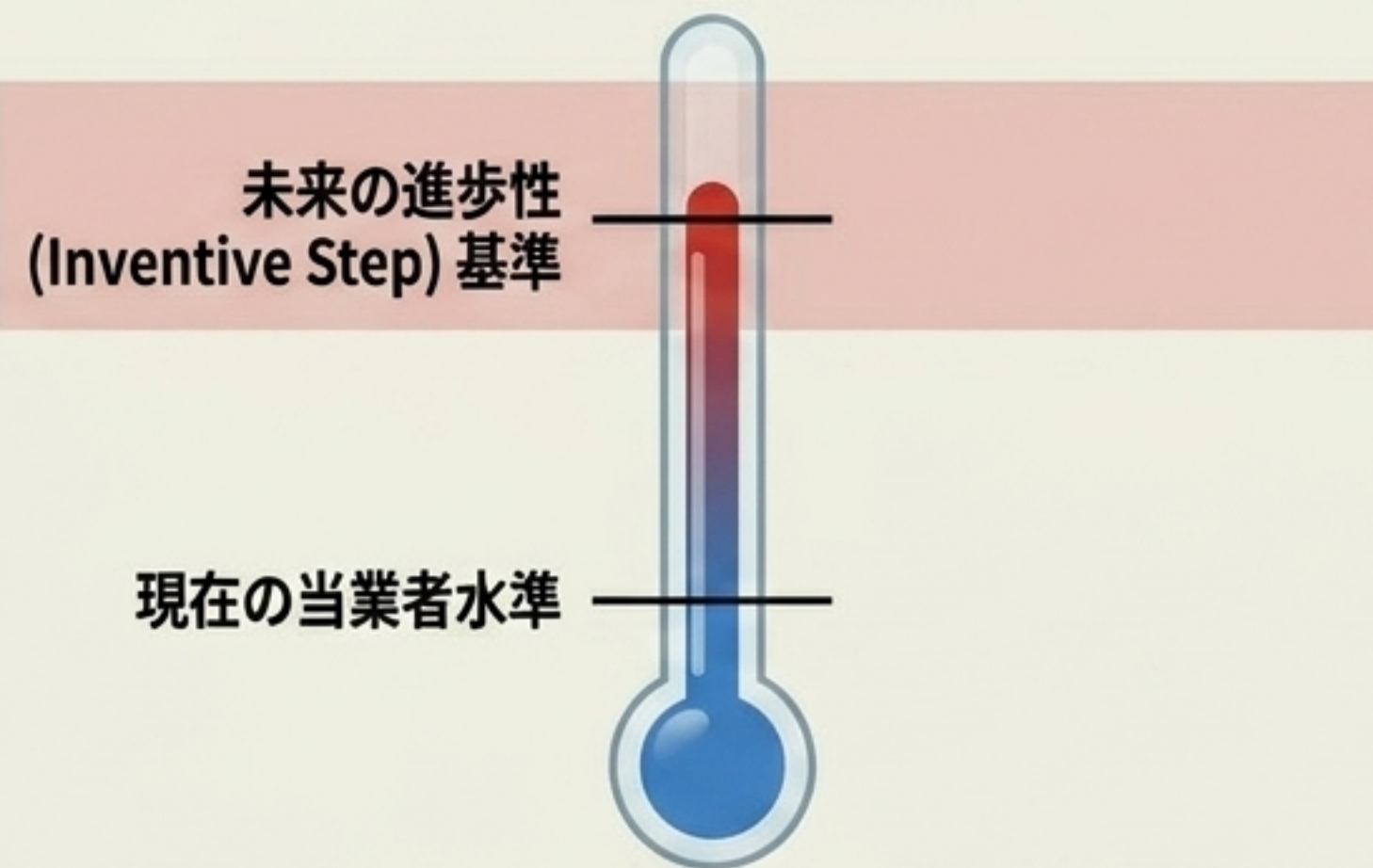
# The 10-Year Threat : 熟練形成の細りと進歩性の壁

## スキルの空洞化 (Skill Hollowing)



AIが形式知を提示し続けることで、次世代が自ら試行錯誤して暗黙知を獲得する経路が絶たれる。

## 進歩性ハードルの上昇



AIツールが「通常的手段」として定着すれば、将来的に特許の進歩性基準そのものが引き上げられ、先行者自身が苦しむ逆説が生じる。



# The Playbook : R&D ・ 知財部門への処方箋

即時 (0～3ヶ月) :  
証跡の紐付け

短期 (3～6ヶ月) :  
データ区別と権限管理

中期 (6～18ヶ月) :  
二層構造の制度化

AIの作業ログ（プロンプト・出力・採否）を実験ノートや発明届と即座に紐付け。人間の「着想」の立証補助資料とする。

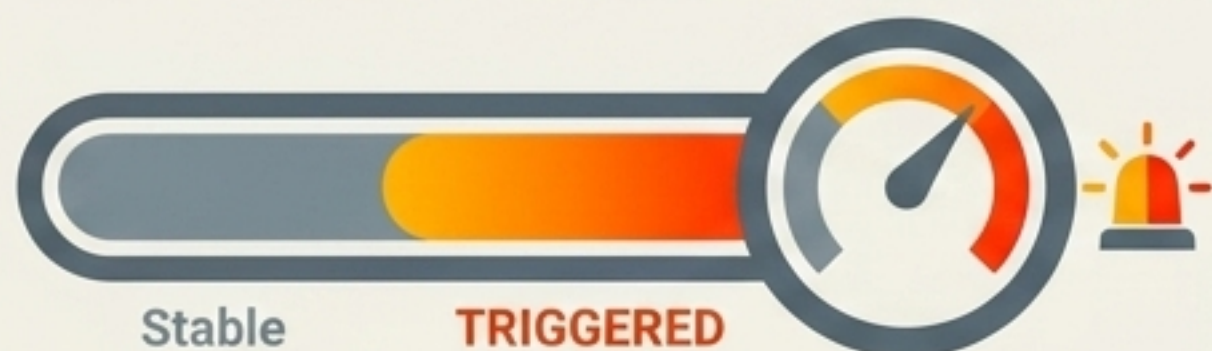
組織共有データについて「営業秘密」と「限定提供データ」を厳格に区別。外部ベンダーとの目的外利用禁止条項を整備。

「仮説はAI生成、検証と最終判断は人間」という二層構造を制度化。利用モデルの記録を必須要件化し再現性を担保。



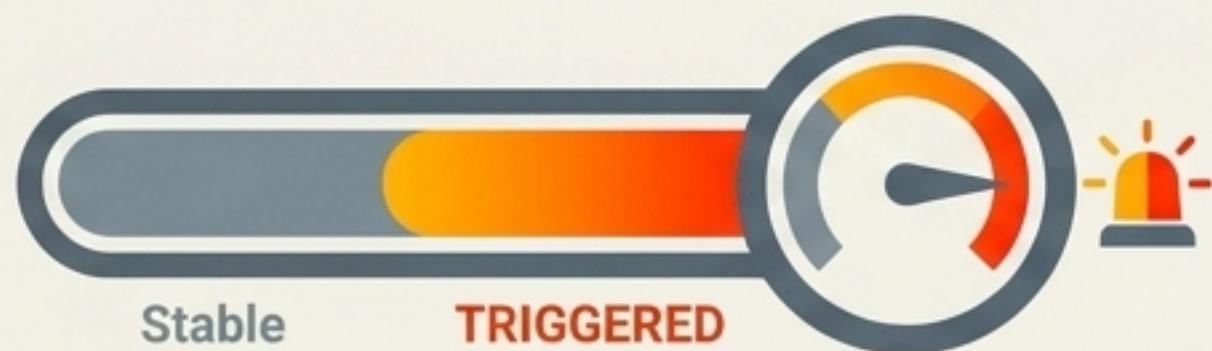
# Dynamic Strategy：見直しのトリガーポイント

以下のシグナルが点灯した瞬間に、知財ポリシーとR&D投資配分を動的に再編する機動力が求められる。



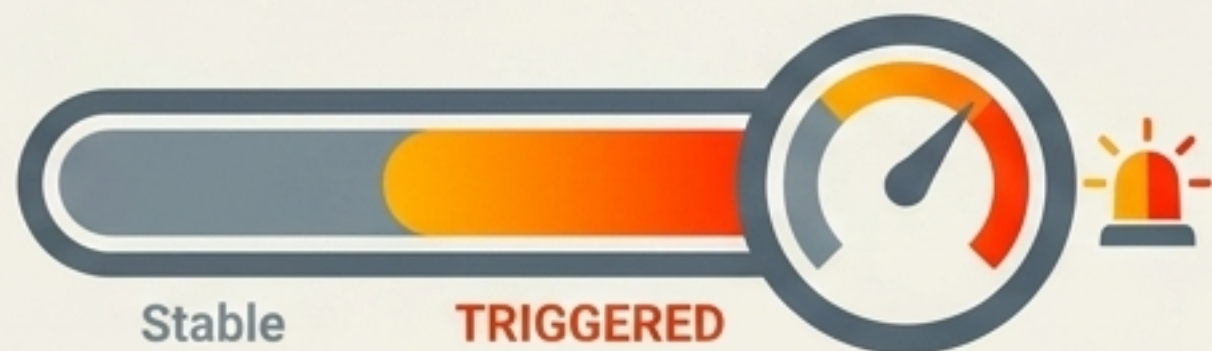
## Trigger 1: 法規制・審査基準の変容

日本の特許法または審査基準が、AI寄与発明の発明者性や「進歩性」について新たな明確な指針を示した時。



## Trigger 2: 紛争リスクの顕在化

AI提案由来の発明・新商品が市場に出現し、帰属や模倣をめぐる自社・他社の法的紛争が発生した時。

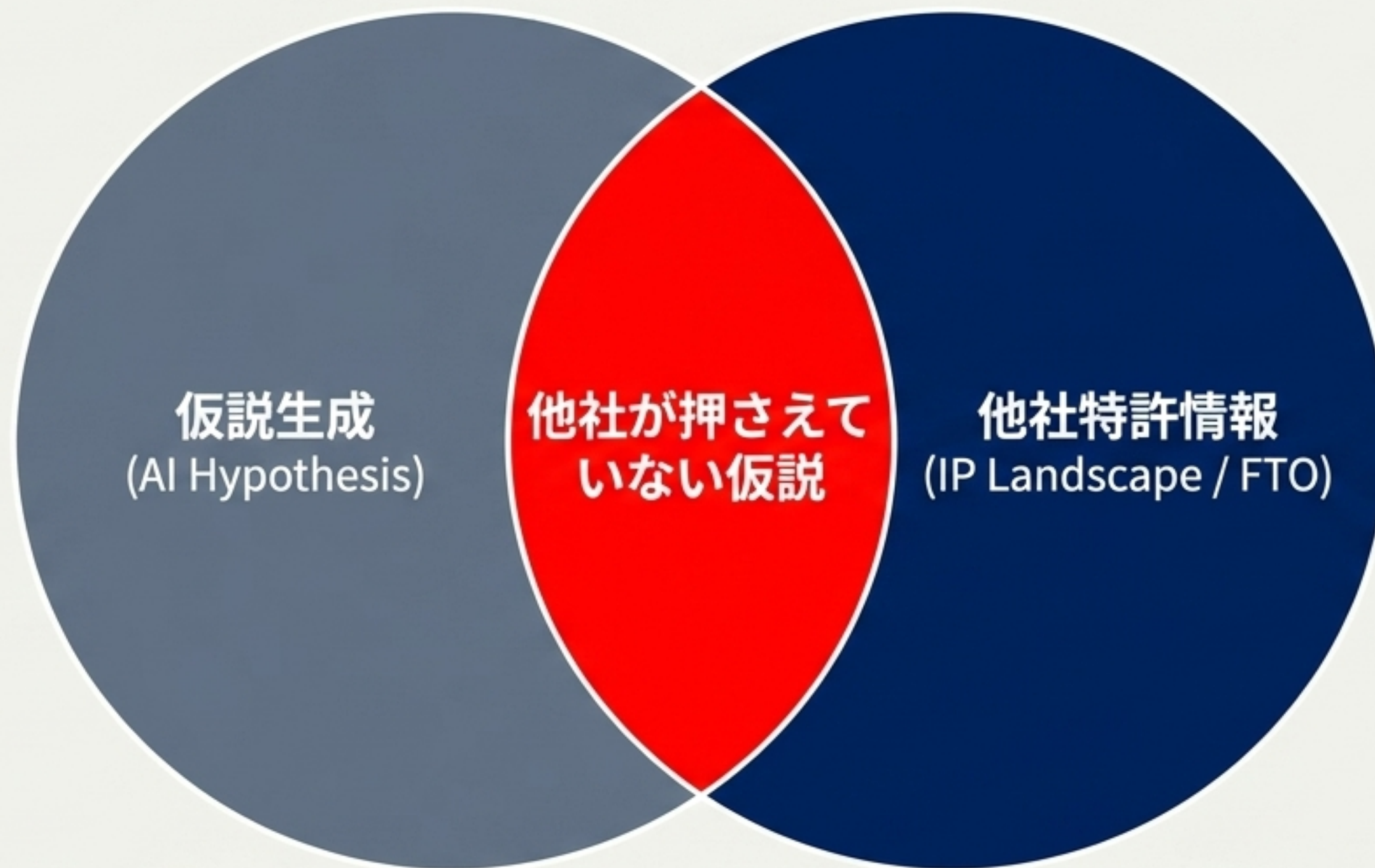


## Trigger 3: 競合の追従と常在型実装

競合他社が効率化目的ではなく、「R&Dプロセスへのエージェント常在型」を本格実装した時。



# The Ultimate Advantage : 究極の競争優位



**最も指摘価値の高い一点:** 単なるAIの常在化は他社も追従する。仮説生成AIと特許情報が同一の探索空間に載って初めて、他社権利を回避した「スイートスポット」だけを優先提示できる。これが下流のボトルネックを一気に解消する究極の一手である。



**「AI研究員」の導入は、  
R&Dデジタル化の終着点ではない。  
知的創造プロセスの再定義の  
始まりである。**

---

ツールとしてのAIに熱狂する段階は終わった。  
問われているのは、激増する「仮説」と冷徹な「特許権」  
の壁をどう接続し、組織の暗黙知を守り抜くかという  
【戦略的知財マネジメント】の構想力である。

