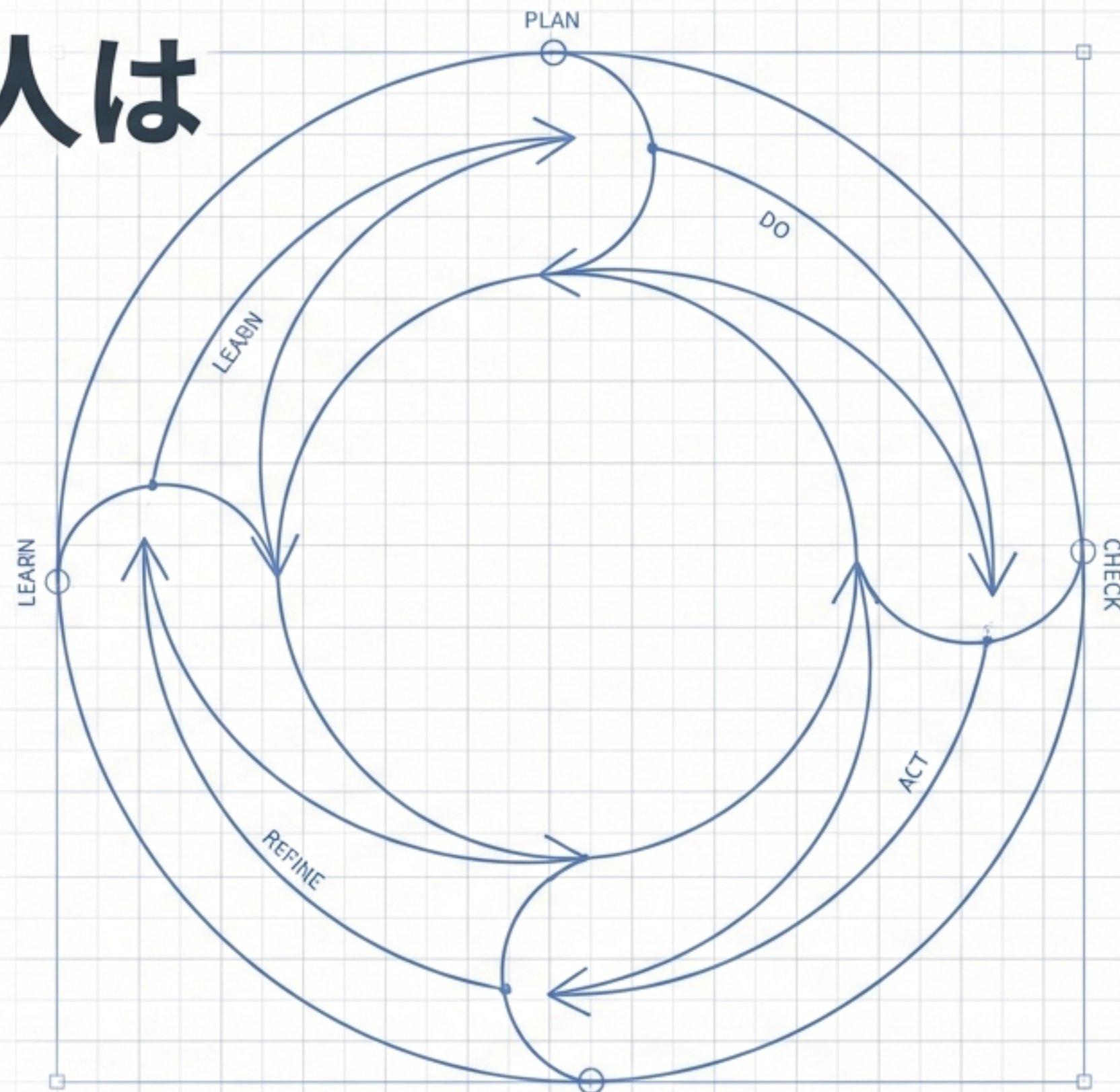


# AI時代に成果を出す人は 「成功」より何を分 析するのか

最終成果から「プロセス能力」への  
KPIシフトと組織実装の青写真

AIの利用率は戦略ではない。  
真の競争優位は  
「検証済み学習速度」にある。



# 単発の「成功」ではなく、失敗確率の低減と学習速度を評価する



## 従来の評価モデル

- **焦点:** 最終成果（売上、利益、リリース数）とAIの「利用量」。
- **誤謬:** 環境要因（運・予算）と個人の実力を混同する。失敗＝悪として隠蔽される。



## AI時代の評価モデル

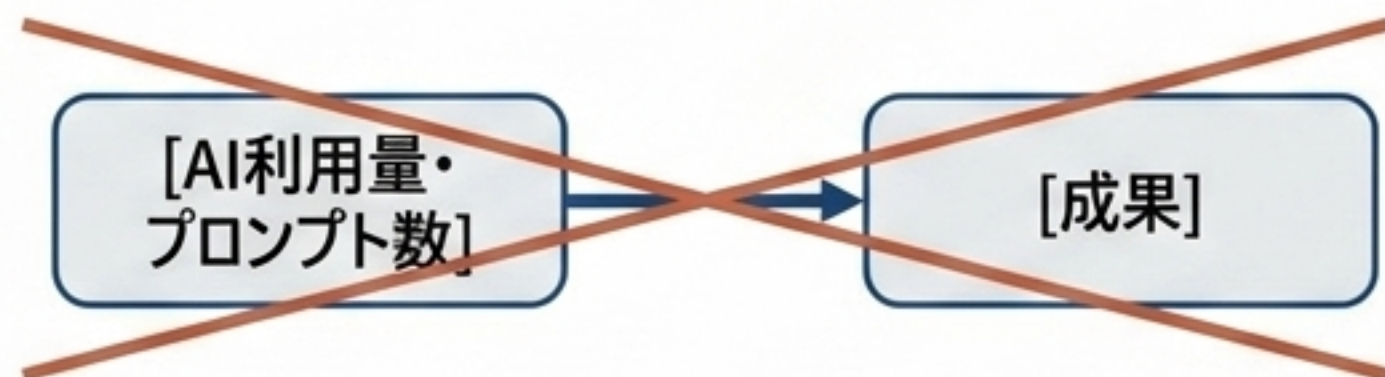
- **焦点:** 失敗確率をどう下げたか。そこからどれだけ速く再利用可能な学習を得たか。
- **真理:** エラーは「情報」である。AI時代における高業績とは、誤差を最速で情報へ変換する能力を指す。

※注意：売上等の「遅行指標」を無視するわけではない。  
診断と改善の焦点を、その手前にある「先行指標」へ移すことが目的である。

# 「AIを何時間使ったか」は虚栄の指標である

BCG/HBSの758人対象実験：AIは能力境界内では速度・品質を劇的に高めるが、境界外(複雑な課題)ではAI利用者の正答率を逆に悪化させる。

## 素朴な直線モデル

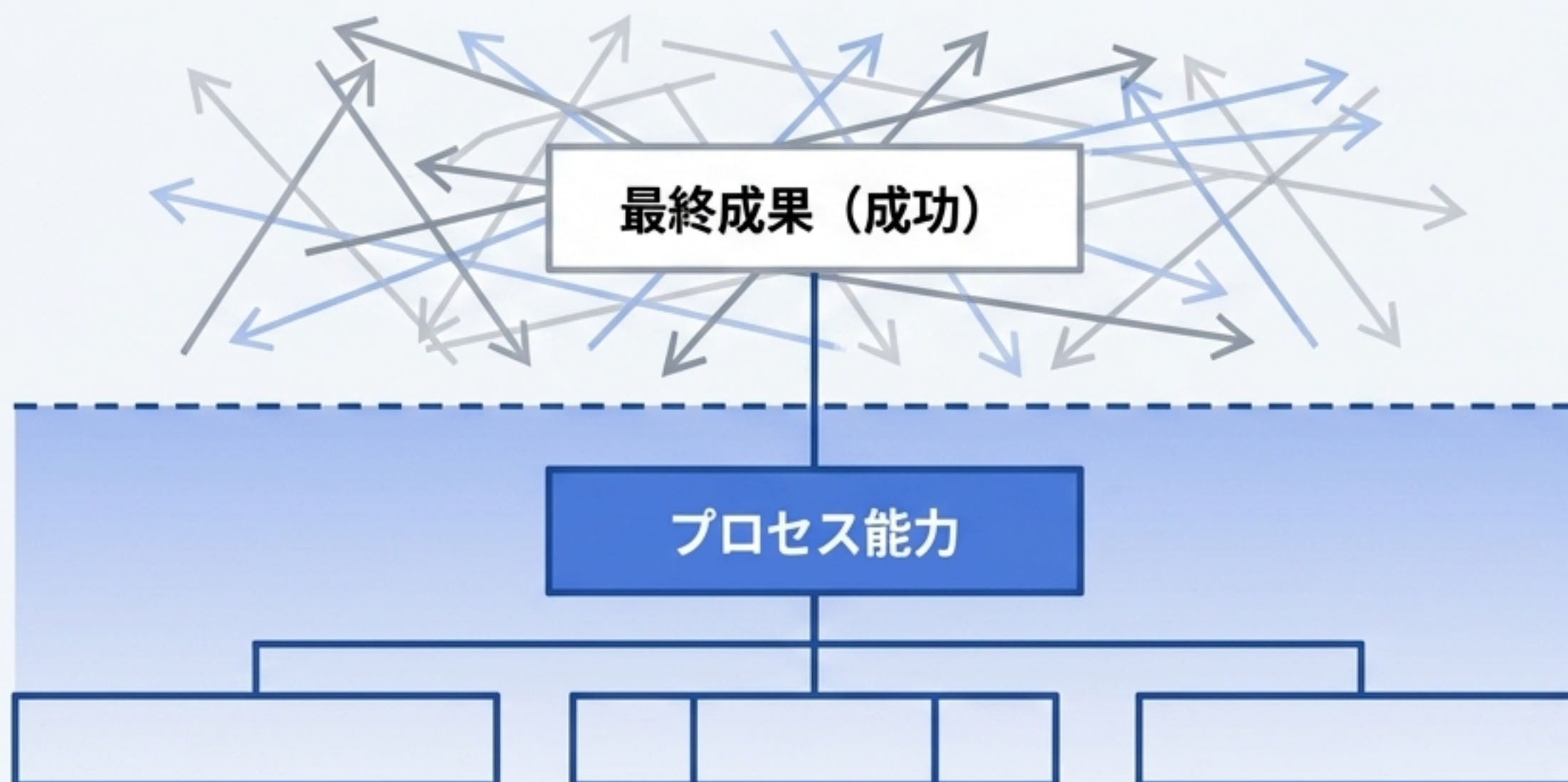


## AI時代の反復学習ループ



価値を生むのは「AIを使ったか」ではなく、「いつAIを疑い、どう修正したか」という境界判断の質である。

# 「成功した人の行動を真似る」ことの危険性

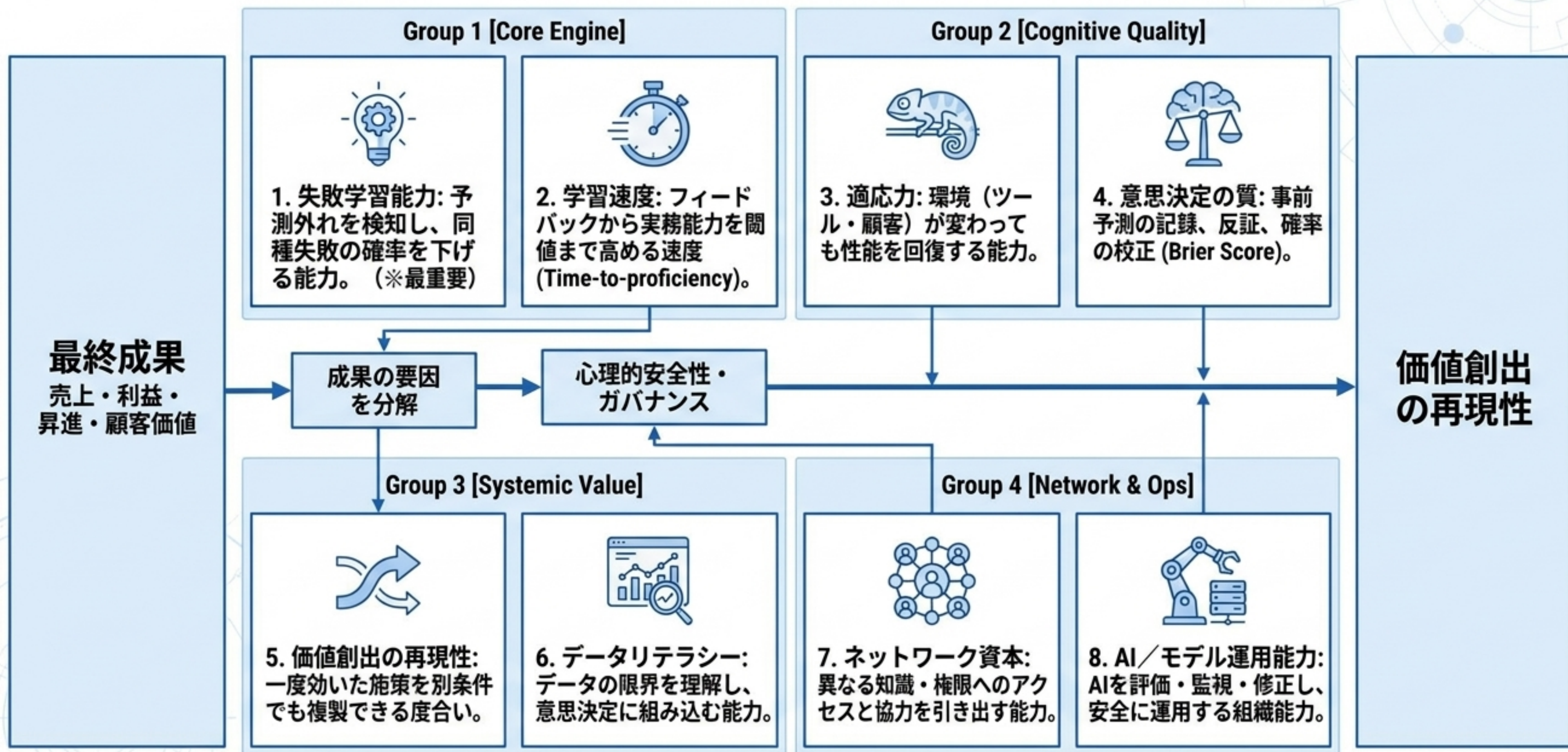


営業受注やプロジェクト成功は、予算、景気、タイミング、競合などの「環境要因（ノイズ）」に大きく左右される。成功事例のコピーは「生存者バイアス」を引き起こす。

Madsen & Desai（宇宙打上げ産業研究）：組織は成功経験より「失敗経験」から効果的に学習し、その知識は減衰しにくい。

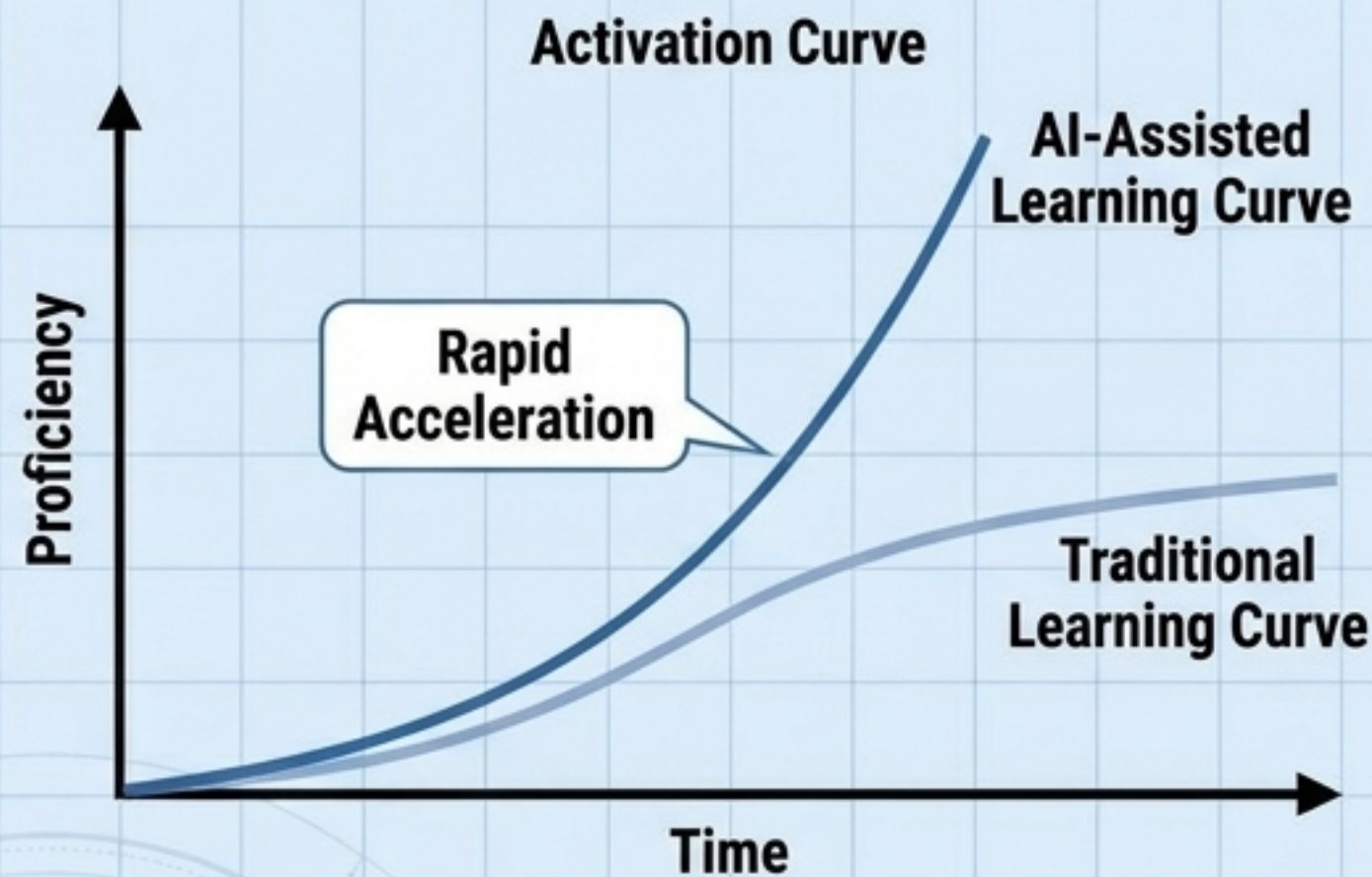
- 従来の問い：「成功したか？」
- AI時代の問い：「何を予測し、何を外し、なぜ外し、どの修正が効き、それが次回も再現したか？」

# 「成功」を説明・予測する8つの先行指標



# 成長のエンジン：失敗学習利回りと学習速度

## 学習速度 (Learning Speed)



Brynjolfssonらの研究：生成AI支援は経験の浅い層の生産性を15%押し上げ、熟練者の暗黙知を伝播させる。

誤ったKPI：研修受講時間（40時間）

正しいKPI：Time-to-proficiency（独力で品質基準を満たすまでの日数）

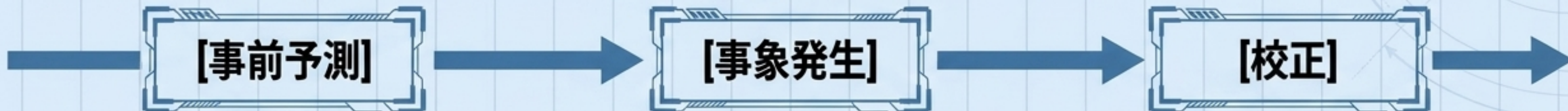
## 失敗学習能力 (Failure Learning Capability)

失敗を減らす目標は「隠蔽」を誘発する。  
重要なのはエラーを情報へ変換する効率。

$$\text{失敗学習利回り} = \frac{\text{分析したインシデント量 (Near miss含む)}}{\div \text{対策後に減少した同種再発・損失}}$$

# 判断の質と価値の再現性をどう測定するか

## 意思決定の質 (Decision Quality)



- 結果バイアスの排除：「結果が良かった＝判断が良かった」ではない。
- 実践：地政学予測トーナメントの知見を応用。事後説明ではなく、「新機能が3ヶ月で利用者を10%増やす確率は70%」と事前登録する。
- KPI：Brier Score（予測確率と実際の結果のズレを測る指標）

## 価値創出の再現性 (Value Reproducibility)

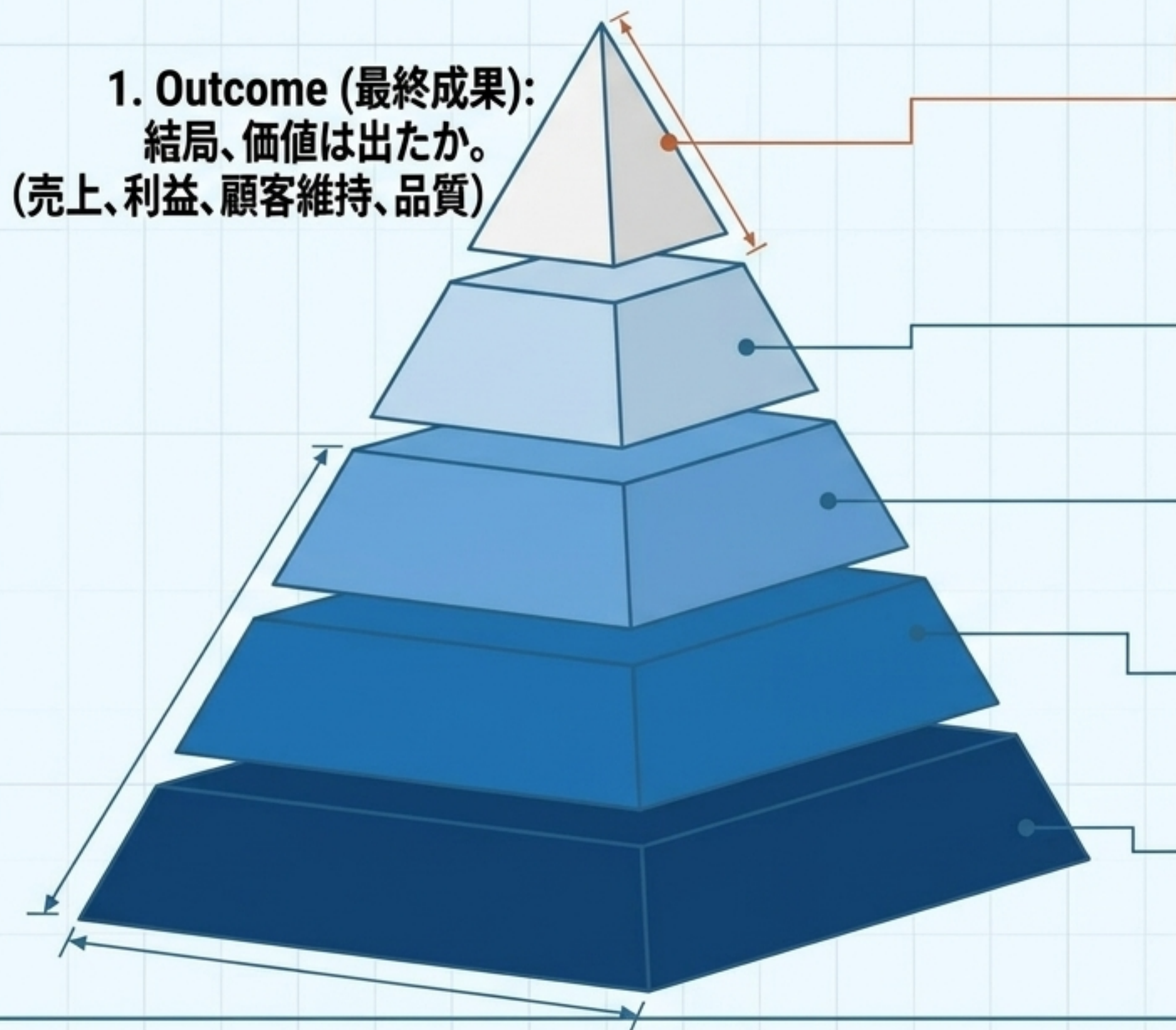
AIによる「一発の大当たり」や「大量のアイデア」は差別化にならない。



# プロセスKPIへの移行を裏付ける科学的エビデンス

| 研究・事例                                  | 主な結果                               | 因果推論の強さ・含意                                                         |
|----------------------------------------|------------------------------------|--------------------------------------------------------------------|
| AI支援 × 学習速度<br>(Brynjolfsson, Science) | 5,172人RCT。<br>低スキル層の学習を促進。         | <div></div><br>[ 証拠: 強 (因果) ]<br>→ 含意: 習熟曲線の短縮を測る。                 |
| AI委譲判断<br>(HBS/BCG)                    | 758人フィールド実験。<br>境界外でのAI依拠は正答率を悪化。  | <div></div><br>[ 証拠: 強 (因果) ]<br>→ 含意: 利用率ではなく「委譲・疑う判断」を評価。        |
| 失敗からの組織学習<br>(Madsen & Desai)          | 宇宙打上げ産業長期観察。<br>失敗経験からの学習は知識減衰が遅い。 | <div></div><br>[ 証拠: 中 (縦断観察) ]<br>→ 含意: 成功DBより失敗・Near Miss DBの構築。 |
| データ駆動意思決定<br>(Brynjolfsson et al.)     | 179社分析。<br>DDD導入企業は生産性が5~6%高い。     | <div></div><br>[ 証拠: 中 ]<br>→ 含意: 閲覧数ではなく意思決定への組み込みを測る。            |
| 心理的安全性 × 学習<br>(Edmondson)             | 51チーム研究。<br>安全性が学習行動を生み成果を媒介。      | <div></div><br>[ 証拠: 関連性 ]<br>→ 含意: 失敗報告の前に発言環境の確保。                |

# 実践的評価フレームワーク：5層構造モデル



最終成果 (Outcome) を消してはならない。  
先行指標だけを追うと「学習したが価値を出していない」状態を正当化してしまう。両者の併記が必須。

2. Learning (学習):  
何を学び、次回の失敗確率を下げたか。  
(再発率、Time-to-proficiency)
3. Judgment (判断): 良いプロセスで決めたか。  
(予測校正、Brier Score、反証数)
4. Leverage (増幅力): 他者・データ・AIを使い価値を増幅したか。(AI受容出力当たり価値、部門横断協働)
5. Guardrail (安全性): 成果のために壊してはいけないものを守ったか。(重大事故、重大誤答率、公平性)

# 職種別：AI時代のKPI適用マトリクス

|                   | 最優先指標                            | 旧来の問題                          | AI時代の適用                                                |
|-------------------|----------------------------------|--------------------------------|--------------------------------------------------------|
| 経営者・<br>事業責任者     | 意思決定品質、資本<br>配分学習速度。             | 売上だけでは市場環境（運）<br>と判断品質を分離できない。 | 大型AI投資前に確率予測・撤退<br>条件を記録し、四半期ごとに<br>判断を校正する。           |
| プロダクト<br>マネージャー   | 検証済み学習速度、<br>ユーザー価値再現性。          | 「リリース数」を増やして<br>も顧客価値は増えない。    | AIでアイデアを大量生成するよ<br>り、重要仮説を反証・確認する<br>までの時間を測る。         |
| エンジニア/<br>MLエンジニア | 失敗学習、MLOps<br>成熟度、AIコード<br>検証能力。 | コード量やモデル単体精度<br>では運用品質が見えない。   | AIによるコード量ではなく、変<br>更後障害の減少と、人手による<br>修正の履歴を評価。         |
| 営業                | 失注学習、予測校正。                       | 担当顧客規模や案件配賦に<br>大きく左右される。      | AIで提案量を増やすのではなく、<br>失注理由別に提案を修正し、再受<br>注率（事前確率の更新）を見る。 |

# ガバナンスと陥りやすい罠（KPIのダークサイド）

## Goodhart化とゲーム化



罠：「AI利用率100%」を目標にすると、不要な仕事にもAIを使う。「再発率ゼロ」は隠蔽を誘発する。



盾：ペア指標の導入。「AI利用回数」には必ず「検証済み価値＋人間修正工数」をセットにする。

## 生存者バイアス



罠：成功者だけを分析すると、成功パターンの過大評価（運の過信）を生む。



盾：成功案件だけでなく、中止案件、失注、モデル棄却、Near Missを必ずデータベースに残す。

## 監視と心理的安全性



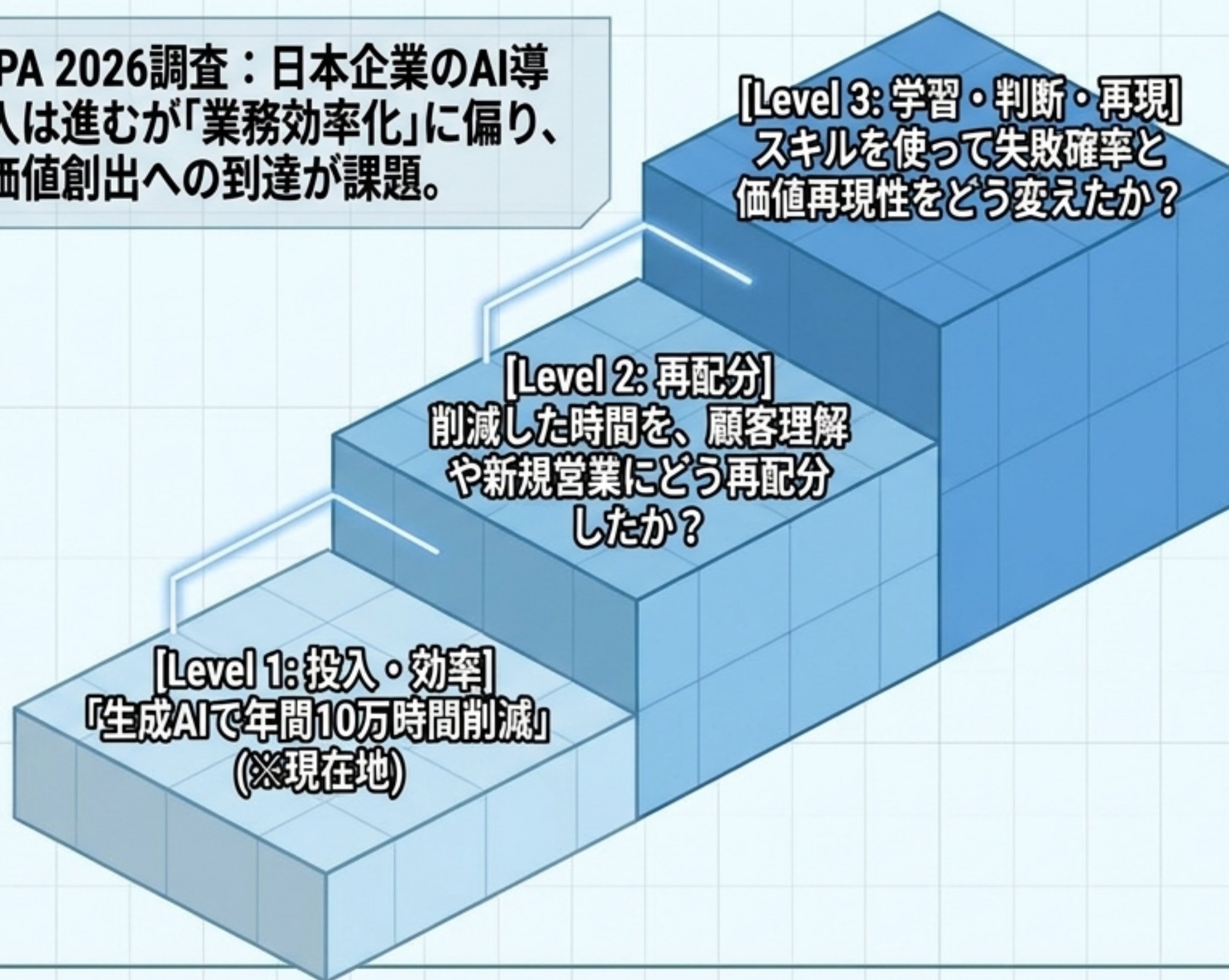
罠：ネットワーク資本を測るためにメールや表情・音声をAIで自動採点し、人事評価に直結させる。



盾：Edmondsonの研究が示す通り、学習行動の前提は「心理的安全性」である。失敗を報告できる環境を壊すような監視利用は厳禁。

# 日本企業における構造的課題の克服

IPA 2026調査：日本企業のAI導入は進むが「業務効率化」に偏り、価値創出への到達が課題。



**Systemic Constraint**  
(METI 2025):

日本の年功的賃金や長期雇用は、個人の学習インセンティブを弱める。

**Actionable Fix:**

スキル獲得 ⇒ 実務での証明 ⇒  
新しい役割へのアクセス ⇒ 処遇  
の連鎖を明示する。個人は  
「ChatGPT Certified」ではなく、  
「どこをAIに委譲し、どこを人  
間が修正したか」という価値の  
ポートフォリオを構築する。

# 新指標の組織実装ロードマップ

## 短期（開始～1ヶ月）

### - 現在地の可視化

- 最終成果を1-2個、先行指標を4-6個選定。
- Decision logと失敗分類の開始。

## 短中期（1～3ヶ月）

### - 指標の妥当性テスト

- 事前確率の記録、ポストモチベーションの開始。
- 職種ごとの基準値（ベースライン）を作成。

## 中期（3～12ヶ月）

### - 因果関係の検証

- 「先行指標の改善」が「成果の改善」につながっているか検証。

## 長期（1年以降）

### - 経営システム化

- 指標と報酬・配置の関係を段階的に連動。
- 組織学習DBの構築。



※絶対のルール:  
いきなりスコアをスコアを給与・  
昇進に結びつけないこと。

- 「先行指標の改善」が「成果の改善」につながっているか検証。
- 効かないKPIは容赦なく廃止する。

# AI時代の真の「トップ人材」とは何か

AIがコード、テキスト、資料作成の限界費用をゼロに近づけるほど、  
「アウトプットの量」の希少性は完全に失われる。

- AIが間違ふ領域を見抜く「判断力」
- 失敗や反証から最速で学ぶ「学習速度」
- 異質な人を巻き込み価値を複製する「再現性」



- AIが間違ふ領域を見抜く「判断力」
- 失敗や反証から最速で学ぶ「学習速度」
- 異質な人を巻き込み価値を複製する「再現性」

AI時代において成果を出し続ける人は、成功回数が多い人ではない。  
自分と組織の誤差（エラー）を最も速く情報へ変換し、次の失敗確率を下げ、  
その知識を他者でも再現可能にできる人である。