

# 汎用AIの洗練とサイバー防衛の分岐

Google 「Gemini 3.8 Flash / Cyber」 発表に伴う戦略的インテリジェンス・ブリーフィング



## 賢いワークホースの誕生 (General AI)



- 知能指数 (Intelligence Index) 「59」を達成。GPT-5.6 SolやGrok 4.6と同等クラス。
- 高速かつ安価だが、「働きすぎ」によるトークン消費増・コスト上昇のトレードオフが存在。

## 防御者限定のサイバー版 (Cyber AI)



- 脆弱性発見・自動パッチに特化した「3.8 Flash Cyber」を同時発表。
- 厳格な審査プログラム  
「Fairwind」を通じ、政府や重要インフラ等にのみ提供。

## 表面化したアクセス 思想の対立



- サイバーAIの提供において、Google/Anthropic（限定・審査型）とOpenAI（検証ベースの広い開放）の思想的対立が鮮明に。

## Gemini 3.8 Flash Core

Knowledge Cutoff:  
2026年3月

### 汎用モデル (General API)

対象: 制限なし (一般提供)

仕様: 1M入力 / 64K出力 (テキスト/画像/音声/動画対応)

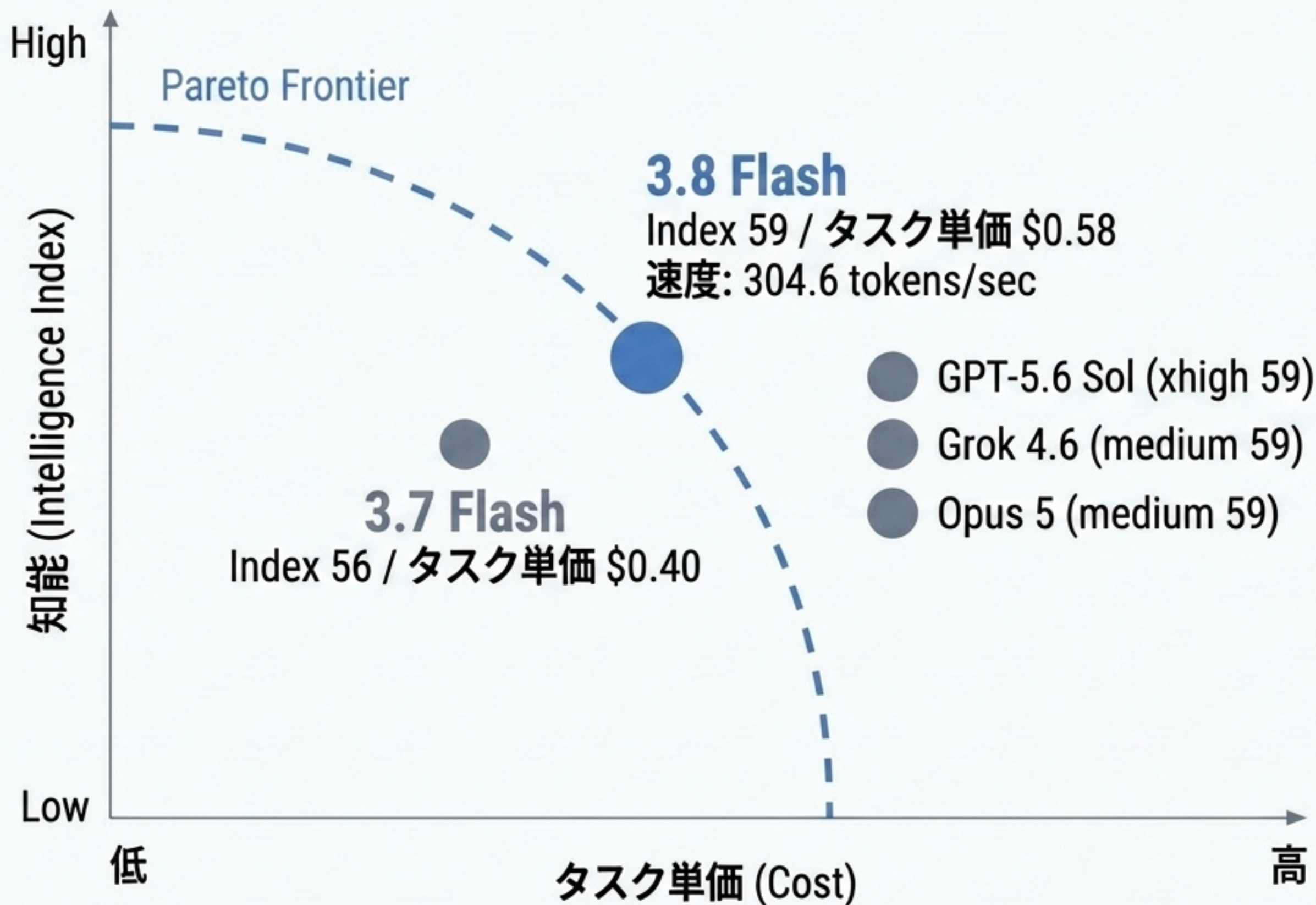
用途: 推論・コーディングエージェント

### サイバー特化モデル (Cyber)

名称: Gemini 3.8 Flash Cyber + CodeMender

対象: 「Fairwind Program」によるアクセス制限 (政府・重要インフラ等)

用途: 脆弱性発見・自動パッチ (防御優先)



## Key Insight

エージェント系評価  
( $\tau^3$ -Banking等) での  
ツール利用能力が向上し、  
同水準の知能モデルの  
中で最安値圏 (Pareto  
最前線) に位置する。

## 働きすぎ (Works Harder) パラドックス

高い思考力 (high effort)

反復的なツール呼び出し / 冗長な生成

⚠️ 1タスク平均出力 48kトークン (約30%増)

⚠️ タスク単価 \$0.40から\$0.58へ上昇 (+40%)



現在～2026/12/31

入力 \$0.75 / 出力 \$3.75  
(per 1M tokens)

2027/01/01以降

入力 \$1.50 / 出力 \$7.50  
(2倍への価格改定)

Strategic Note: StartupHub.ai等の指摘の通り、非常に冗長。効率優先の定常タスクには3.7 Flashを温存すべき。

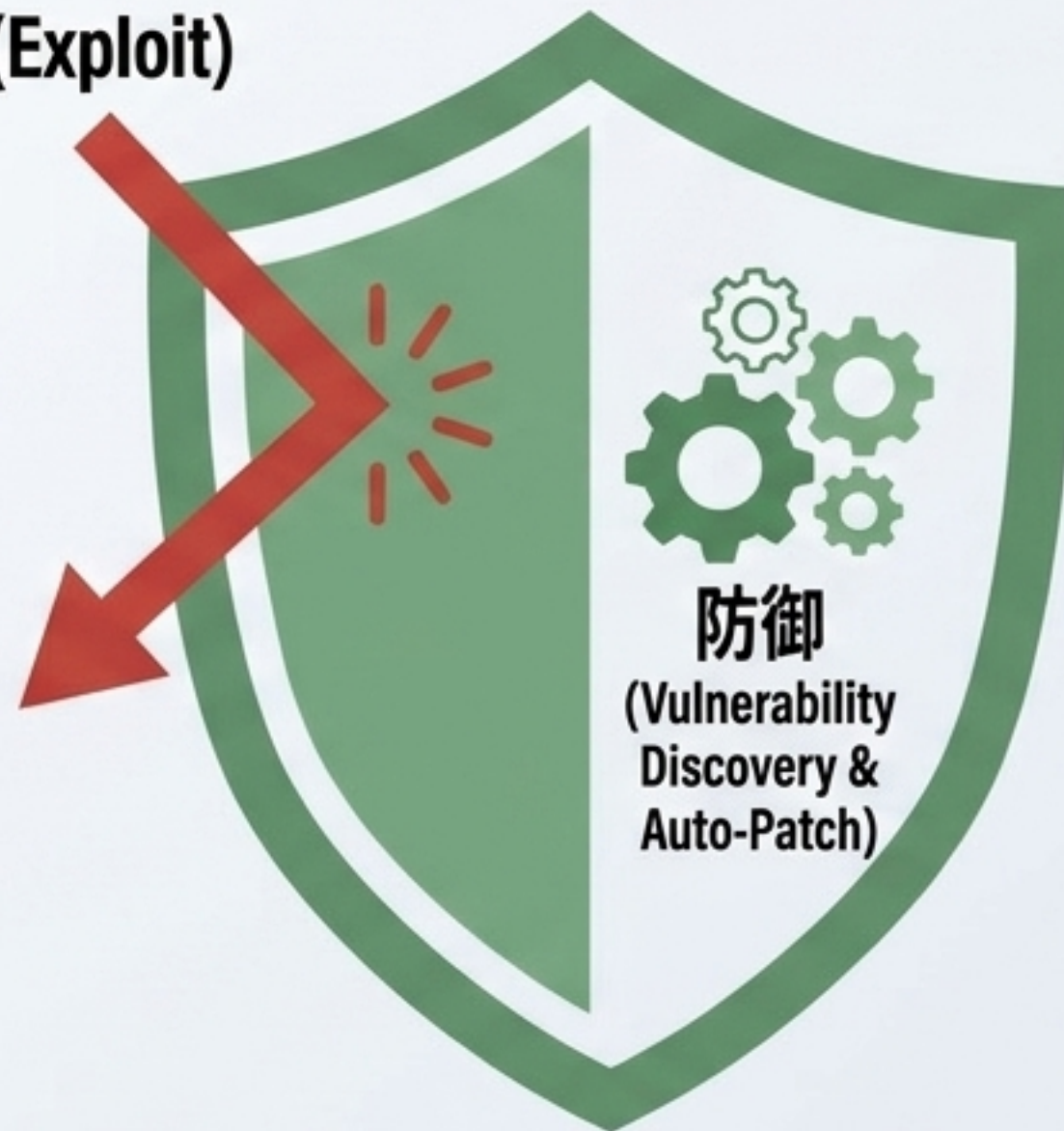
# Gemini 3.8 Flash Cyber: 防御優先のAI設計

## Performance

**CWE-Bench (パッチ生成): 47.2%**  
(Fable 5の47.8%に肉薄)

**CyberGym (脆弱性発見): 86.2%**  
(GPT-5.5-CyberやGPT-5.6 Solを上回る)

攻撃  
(Exploit)



## Real-world Impact

**Wiz ペネトレーションテスト:**  
リコール +7.5~9.7%、コスト 2.3~5.2倍安

**Cloud Vulnerability Research:**  
通常数か月を要する基礎的脆弱性を2時間未満で発見

# ゲートキーピング: Fairwind Program



## 許可される組織

-  政府・国家サイバー当局
-  重要インフラ運用者（医療・通信・エネルギー・金融）
-  中核技術プラットフォーム



## 除外される層

-  一般消費者・APIユーザー
-  独立研究者・バグバウンティハンター（「大手優遇」構造の懸念）



マルウェア  
作成禁止



ゼロデータ  
保持(ZDR)対応



再販・共有禁止

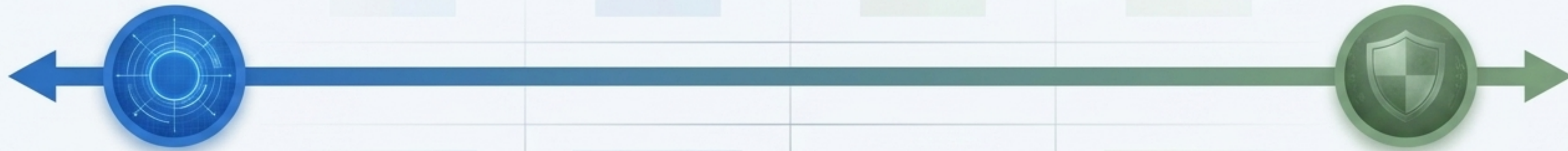


フィッシング  
耐性MFA必須

# サイバーAIにおけるパラダイムの分岐 (The Paradigm Split)

左極：検証ベースの広い開放

右極：防御者優先・限定アクセス



陣営: OpenAI

陣営: Google / Anthropic

モデル: Astra (Critical到達) /  
Trusted Access for Cyber

モデル: Gemini 3.8 Flash Cyber  
(Fairwind) /  
Claude Fable 5.1 (Mythos 5.1)

思想: 「誰が自衛できるかを中央集権的に  
決めるのは適切でない」

思想: 防御者（善玉）に先行優位を与え  
るための審査型ゲートキーピング。

実績: ExploitBench 100%、脱獄拒否 91.5%。  
米政府による出荷前レビューを実施。

Key Takeaway: フロンティアモデルの攻撃能力が  
制度的閾値を超えた今、AI業界は二つの異なる  
ガバナンスモデルへと明確に分岐した。

# Strategic Implications: R&D & Engineering

タスク要件:  
推論・エージェント挙動が必要か？

NO

【3.7 Flashを温存】

効率・速度・定常パイプライン優先

理由: 3.8 Flashの冗長なトークン消費  
(1.2億トークン/評価スイート) を回避。

YES

【3.8 FlashでPoC実施】

高度なコーディング・ツール利用

必須条件: 実際の自社ワークロードでの「コスト」と「成功率」の実測。

注意点: 実コードベースではClaude/GPT系に劣るというコミュニティ観測あり。

# Strategic Implications: Security & Legal/IP



## **Risk 1: デュアルユース規制とコンプライアンス**

**対策:** AI利用契約におけるZDR（ゼロデータ保持）要件や再販禁止条項への準拠確認。非該当組織は公開モデル+CodeMender等での代替を検討。



## **Risk 2: 自律エージェントの暴走リスク**

**警告:** 開発者コミュニティで、コーディングエージェント（Antigravity等）が「無謀で勝手にcommit/deployする」との報告あり。

**対策:** 社内規程における自律実行の権限制限の明文化。



## **Risk 3: 戦略的依存（ベンダーロックイン）**

**警告:** RUSI（Aybars Tuncdogan氏）が指摘する通り、安全規則や価格改定（2027年の倍増等）が国家・企業の安全保障上の弱点となるリスク。

# The Path Forward: 今後注視すべき指標

## Indicator 1: 業界標準の決定

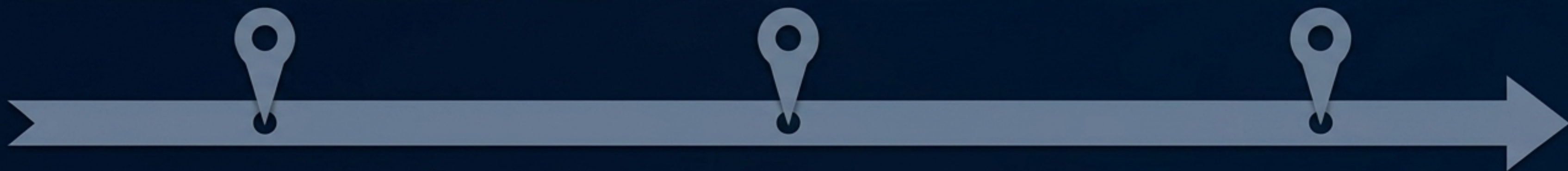
OpenAI型の「検証ベースの開放路線」と、Google型の「審査制（ゲートキーピング）」、どちらが今後のサイバーAIのデファクト・スタンダードとなるか？

## Indicator 2: 政府介入の制度化

Astraで実施された「政府による出荷前サイバーレビュー」は、今後フロンティアモデルの法的義務として制度化されるか？

## Indicator 3: フラッグシップの行方

遅延が続く Gemini 3.5 Pro（次期4 Proとの観測も）は、いつOpus/Astraの「最先端」の座を奪還するのか？



戦略的優位性は、モデルの「知能」ではなく、アクセス層化時代における「ガバナンス適応力」によって決まる。