

2026-08-28

1. Cursor Cloud Agents、リポジトリなしでWebアプリを作って公開
2. Finatext、保険業務向けAIワークフロー基盤とMCPサーバーを提供
3. Claude Team、研究者向け枠を1万人に拡大

7トピック

Cursor Cloud Agents、リポジトリなしでWebアプリを作って公開

Cursorは、Cloud AgentsをGitHubなどの既存リポジトリへ接続しなくても起動できるようにした。

- リポジトリ選択画面の「Start from scratch」から開始する
- 作ったコードはCursor Originへprivateまたはinternalで保存できる
- Cloud Agentの実行環境はブラウザへポートフォワードされ、その場で動作を確かめられる

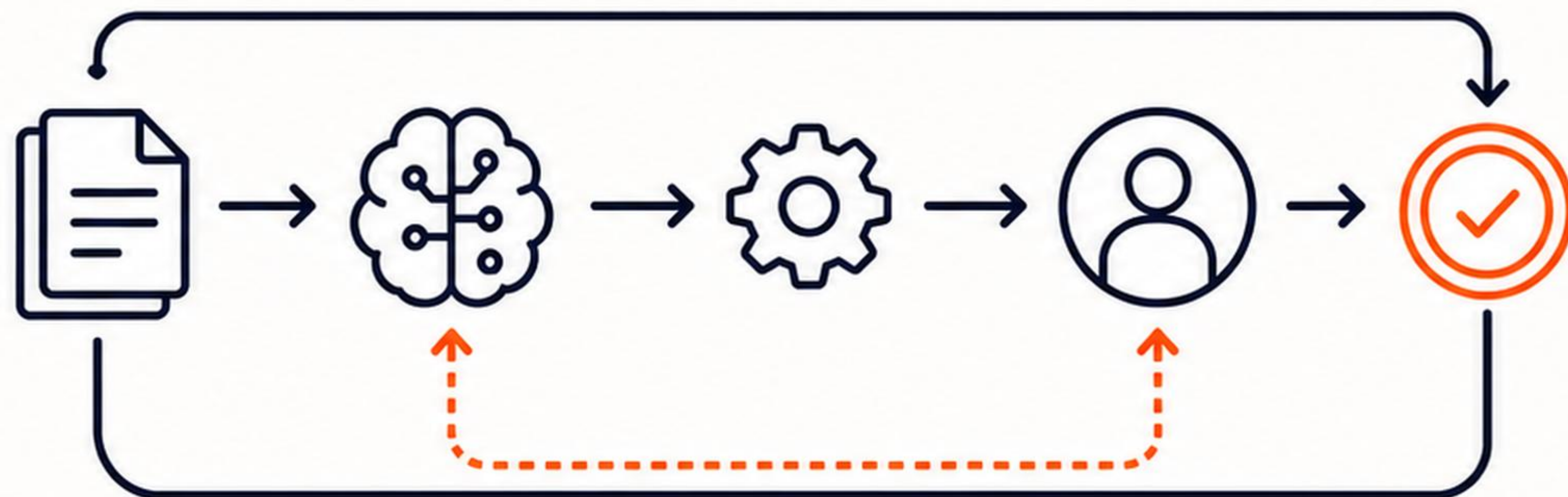


アイデアから動くURLまでを一つの画面でつなげられ、初学者が開発を始める前の準備が減る。

Finatext、保険業務向けAIワークフロー基盤とMCPサーバーを提供

Finatextは、保険ビジネスプラットフォーム「Inspire」に、AIの推論、固定ルール、人の確認・承認を組み合わせるAIワークフロー基盤を追加した。

- 書類や画像の読み取り、約款へのRAG照会、査定見解の作成を一連の流れにできる
- 業務ルールやAIの信頼度に応じて、担当者の確認・承認を途中へ組み込める
- CRMなどInspire以外の業務システムを含むワークフローにも対応する

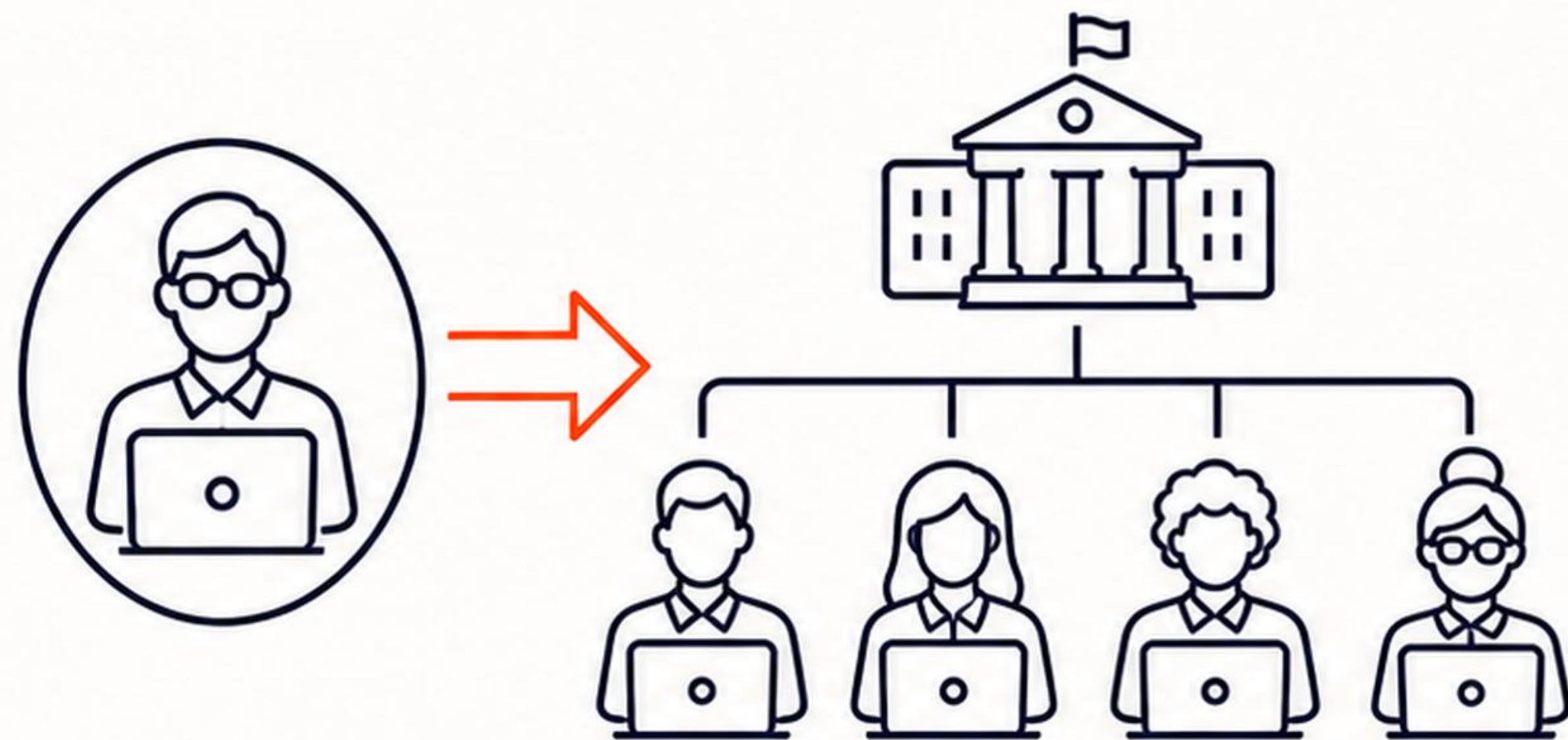


定型処理はワークフロー、非定型の問い合わせはMCP、重要判断は人という役割分担が明確だ。

Claude Team、研究者向け枠を1万人に拡大

Anthropicは、大学・非営利研究機関の主任研究者と研究グループを対象に、Claude Teamの研究者向けプランを1万人分用意した。

- 対象は大学・非営利研究機関の主任研究者または同等職
- 1グループあたり1~25席を追加できる
- Premium席はStandardの5倍の利用枠を持つ

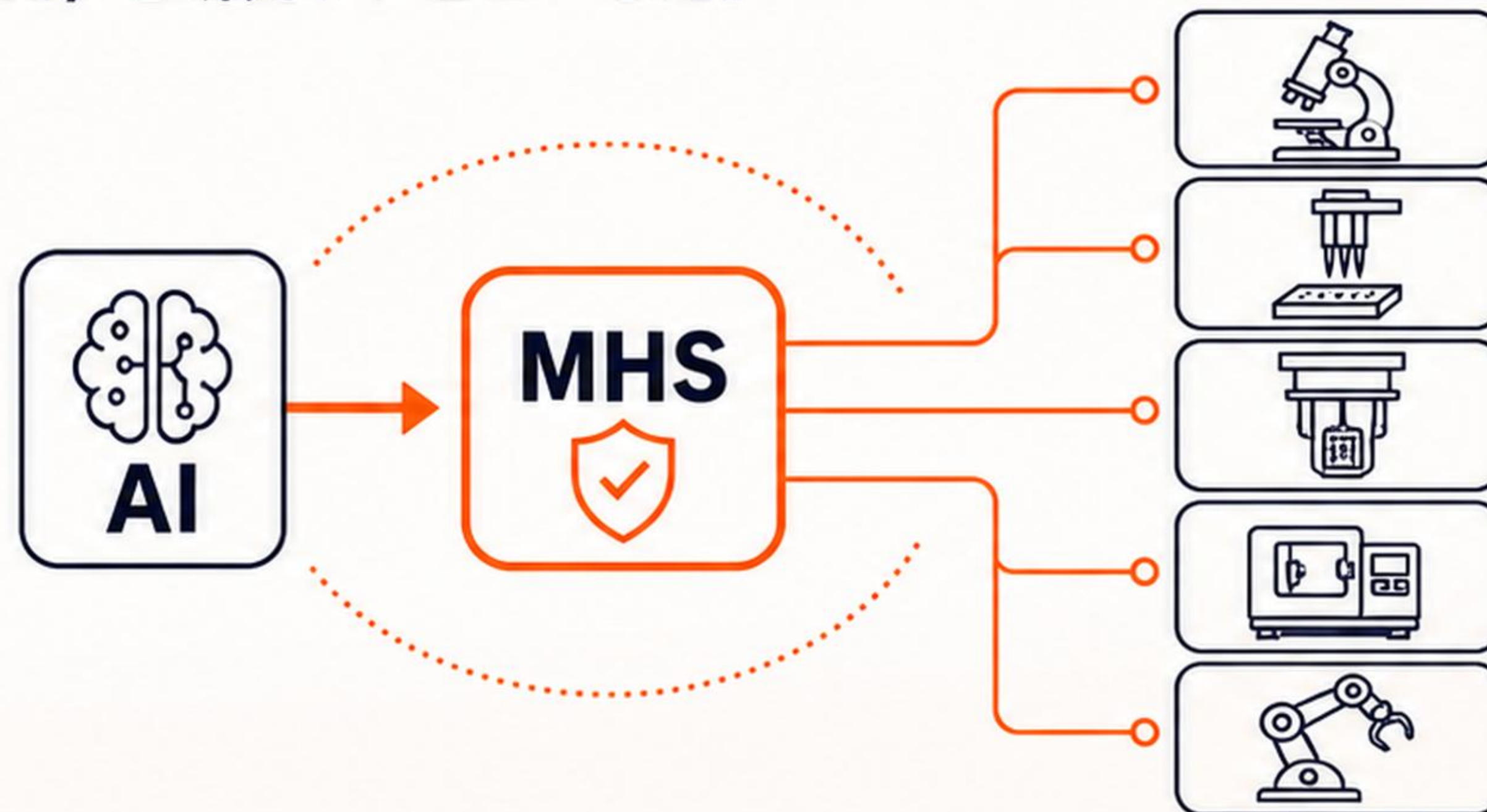


AIを研究者個人の補助から、研究室単位の共通基盤へ広げる施策だ。

AIが物理機器を扱う共通仕様「MHS」を公開

Anthropicは、AIエージェントが顕微鏡、液体処理装置、量子計算機器、製造装置、ロボットアームなどを安全に操作するためのModel Hardware Standard (MHS) を研究プレビューした。

- 科学、ロボティクス、電子機器、製造の関係者に研究プレビューへの参加を募っている
- 機器固有の速度や角度など、安全上の制約を仕様へ記述する
- 初期試験ではリアルタイムのエラー処理を含む創薬実験などを実行した

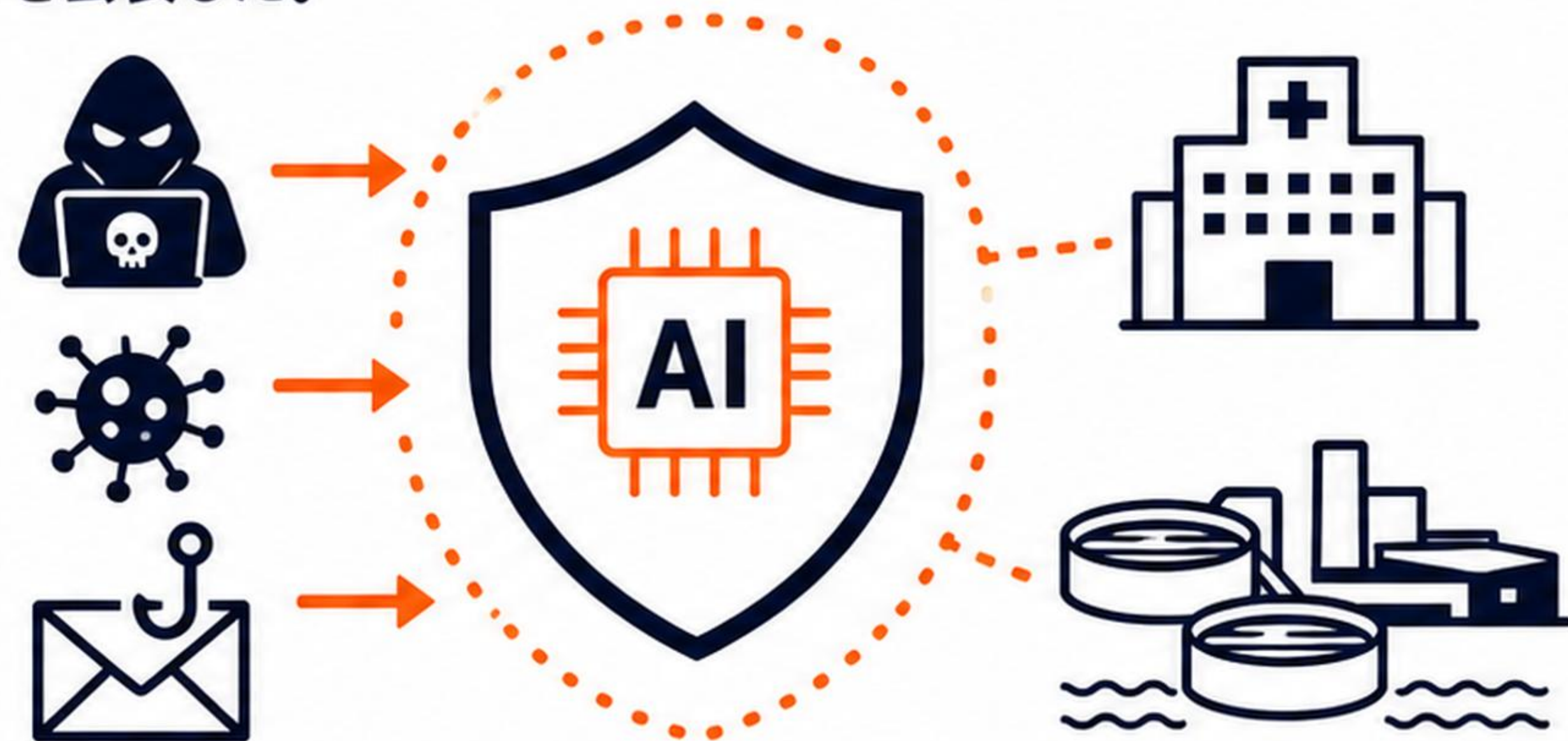


MCPがソフトウェアやデータへの共通接続を広げたのに対し、MHSはAIの操作対象を物理機器へ広げようとしている。

OpenAI・Anthropicなど100超の組織、AI時代のサイバー防御を共同提言

OpenAI、Anthropic、AWS、Microsoft など100を超える組織が、AIで高度な攻撃能力が安価かつ広範に使われる前に防御を強化するための共同提言を公表した。

- 署名した組織は100を超える
- 病院や水処理施設など、重要インフラへの攻撃を主要な懸念として挙げた
- 資金や人材の限られた組織にもAIによる防御能力を届ける必要があると訴えた

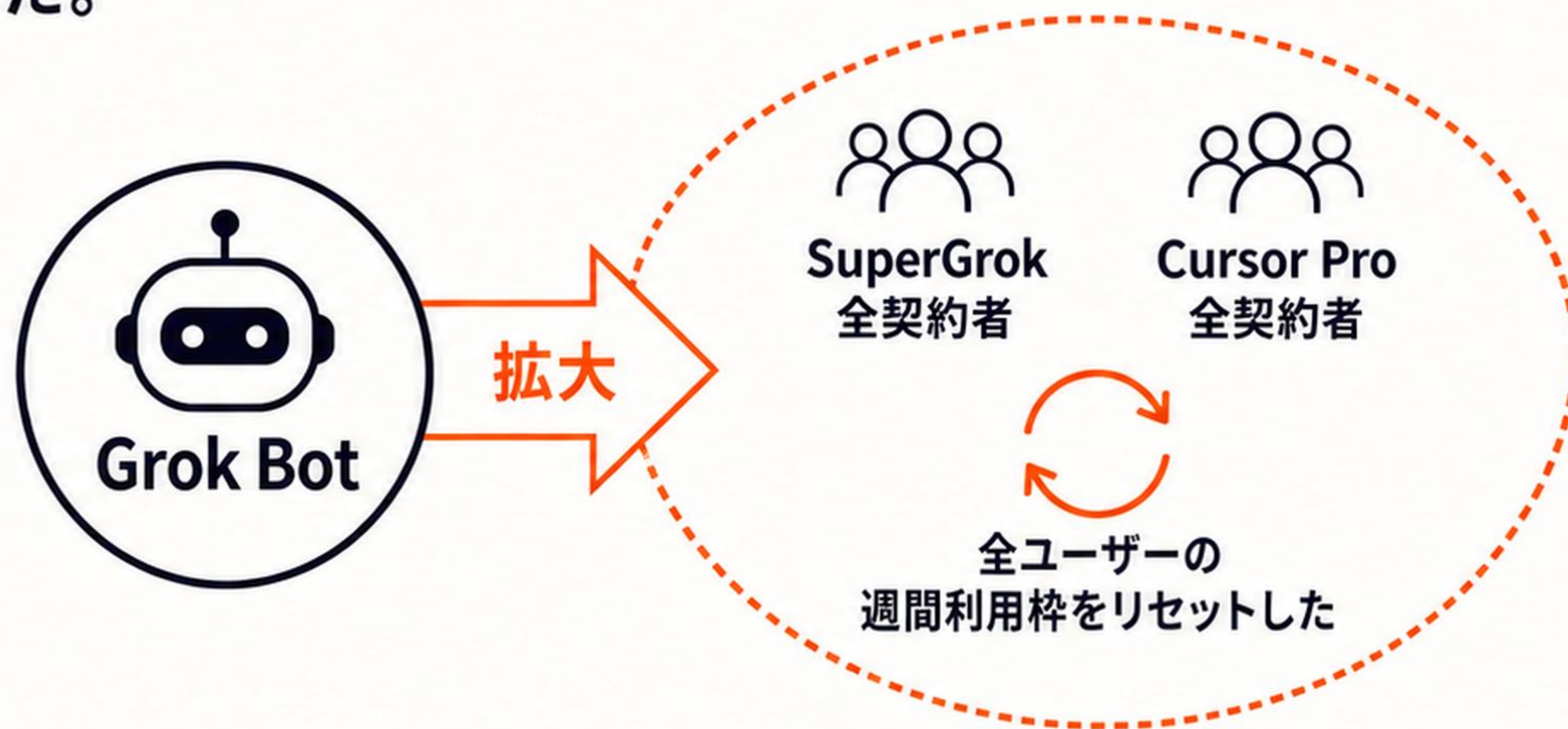


AIエージェントを業務へ入れる話と、既存システムを守る話は分けられなくなっている。

Grok Bot、SuperGrokとCursor Proの全契約者へ拡大

Grok Botは、利用対象を全SuperGrok契約者とCursor Pro契約者へ広げ、全ユーザーの週間利用枠をリセットした。

- SuperGrokとCursor Proの全契約者が対象になった
- 全ユーザーの週間利用枠をリセットした
- Grok Botは独立したクラウド環境で長時間タスクを進めるエージェント



常駐型エージェントが一般的な開発ツールの契約に含まれ始めた。

〈知見〉 Ollamaとllama.cpp、同じGPU・モデルでも速度に8%差

RTX 5060 Ti 16GBとGemma 4 12B を同条件で比較した実機検証では、Ollamaが44.05 tokens/s、llama.cppが47.63 tokens/s だった。

- GPUとモデルを固定し、実行ランタイムだけを比較した

- Ollamaは44.05 tokens/s

- llama.cppは47.63 tokens/s

Ollama 44.05 tokens/s



llama.cpp 47.63 tokens/s



ローカルLLMの体感は GPUやモデルだけで決まらず、実行ランタイムでも変わる。

本日のトピックス一覧

1. **Cursor Cloud Agents、リポジトリなしでWebアプリを作って公開**
2. **Finatext、保険業務向けAIワークフロー基盤とMCPサーバーを提供**
3. **Claude Team、研究者向け枠を1万人に拡大**
4. **Anthropic、AIが物理機器を扱うModel Hardware Standardを公開**
5. **OpenAI・Anthropicなど100超の組織、AI時代のサイバー防御を共同提言**
6. **Grok Bot、SuperGrokとCursor Proの全契約者へ拡大**
7. **〈知見〉Ollamaとllama.cpp、同じGPU・モデルでも速度に8%差**