

Discovery Loop

ジェフ・ディーン退職の意味、AI 研究自動化、 再帰的自己改善（RSI）への距離

GPT-5.6 Sol

中心命題

これは Google との全面決別ではなく、創業者級の人材を独立組織へ移しつつ、Alphabet が資本・計算資源・共同研究で接続を残す「協調型独立」である。技術的には強い RSI の実現ではなく、その前段となる ML 研究の閉ループ自動化への大型ベットだ。

調査基準日 2026 年 8 月 15 日（日本時間）

分析範囲 事実確認 | 技術 | 競争 | Google | 資金・事業 | 安全 | 知財 | 日本への示唆

注意 公開情報に基づく独立分析。確率は独自推定であり、投資助言ではない。

エグゼクティブ・サマリー

BOTTOM LINE

設立・4人の共同創業・Alphabet との関係は確認済み。一方、製品・論文・顧客・ベンチマークは未公表で、10 億ドル調達／100 億ドル評価は交渉中との報道にすぎない。2026 年 8 月時点の企業価値は、実績よりも創業者の信頼性と将来の実験ループに置かれている。

6 つの結論

- ① **事実の骨格は正しい**。2026 年 8 月 5 日（米国時間）、Jeff Dean、Sanjay Ghemawat、Quoc V. Le、Oriol Vinyals が Discovery Loop の設立と Dean の退職予定を発表し、Dean は翌 6 日に Google を退社した。同社はデラウェア州 PBC である。Dean の Google 在籍は 27 年だが、「27 年間 Chief Scientist」だったわけではない。[\[1\]\[2\]\[4\]](#)
- ② **RSI は言い過ぎに注意**。公式ミッションは、実験の提案・実行・評価・学習を自動化し、まず ML 研究で自社スタックを改善すること。これは RSI へつながり、主要投資家も再帰的自己改善として位置づけるが、完全自律で後継基盤モデルを連続生成する「強い RSI」を会社が公式に公約・実証した事実はない。[\[2\]\[3\]\[6\]\[11\]](#)
- ③ **最大の強みはフルスタック**。4 人はモデルだけでなく、分散システム、TPU、学習基盤、評価、エージェント、製品化までを横断する。研究アイデアを数千の計算実験へ変換するには、この統合能力がモデル性能以上に効く。[\[2\]\[3\]](#)
- ④ **Google にとって損失であり、同時にヘッジ**。希少な暗黙知と採用力を失う一方、Alphabet は創業投資家、Cloud/compute partner、共同研究相手として関係を残す。内部の Gemini 製品化と、外部の長期探索を分業する構図にも見える。[\[1\]\[4\]\[7\]\[8\]](#)
- ⑤ **最有力の初期成果は「研究 OS」**。新しい万能モデルより先に、学習カーネル、分散処理、データ配合、ポストトレーニング、評価、実験スケジューリングを自動化し、検証済み改善 1 件あたりの時間・費用を下げる可能性が高い。
- ⑥ **2030 年の本命は Human-on-the-loop**。人間が課題、制約、評価器、最終採否を管理し、AI が実験実務の大半を担う形が中心シナリオ。強い RSI の 2030 年までの実現は上振れであり、現時点で基本ケースに置くべきではない。

情報の読み分け

区分	本レポートでの意味	代表例
確認済み	会社・当事者・法務顧問・一次資料で確認	設立、4 人の共同創業、PBC、Alphabet との提携
報道段階	関係者等に基づくが未確定	10 億ドル調達・約 100 億ドル評価の協議
未確認／注意	公式証拠がなく、既成事実化すべきでない	製品稼働、数千実験の実績、Google IP の移転、完全 RSI
分析・予測	公開情報からの推論。確率は独自推定	研究 OS 化、事業モデル、2026–2030 年シナリオ

目次

章	テーマ	問い
1	現状とファクトチェック	何が確定し、何がまだ物語か
2	戦略と技術アーキテクチャ	どこから自動化し、何が複利化するか
3	RSI への距離	現在の自己改善と強い RSI は何が違うか
4	競争・技術ボトルネック	誰と競い、何が失敗要因になるか
5–8	Google、資金、安全、知財	産業構造と制度はどう変わるか
9–11	予測・日本への示唆・総合判断	何をいつ見て、どう動くべきか

1 | 何が起きたのか：現状とファクトチェック

Discovery Loop の話題は、実在する会社設立と、まだ検証されていない将来像が同時に流通している。最初に両者を分離する必要がある。

1.1 2026 年 8 月 15 日時点の事実台帳

判定	主張	根拠・留保
確認済み	米国時間 8 月 5 日に 4 人が設立を発表	会社、Google、法務顧問が一致。Dean は翌 6 日に Google を退社。Discovery Loop, Inc.はデラウェア州 PBC。[1][2][4][5]
確認済み	目的は実験ループ全体の自動化	まず ML 研究・エンジニアリング。自社を最初の顧客とし、その後、科学・工学へ。[2]
確認済み	Alphabet が創業投資家・Cloud／compute partner	ML systems と関連インフラの研究枠組みでも協力。金銭条件・独占性は非公開。[1][4]
確認済み	初期ラウンドを Radical と Khosla が共同主導	Lightspeed、Kleiner Perkins、Doerr Capital、Alphabet が参加。8 月 9 日時点でクロージング途上。[3][4][5]
報道段階	10 億ドル調達・約 100 億ドル評価	8 月 13 日の関係者報道。条件は変わり得て、会社はコメントせず。成立済みと扱わない。[9]
未確認	稼働製品、顧客、公開ベンチマーク	公式サイトはミッション、チーム、求人が中心。公開技術成果はまだない。[2][10]
注意	「RSI を公式に掲げ、実現へ」	会社公式は RSI を看板語にしていない。ML 研究の自己改善ループは RSI に接続するが、強い RSI とは別。[2][6][11]

1.2 創業チームが持つ「結合された能力」

重要なのは著名人が 4 人いることより、Google で別々だった層を 1 つの創業チームがつかないでいることだ。以下の実績は各人の単独発明ではなく、本人が主導・貢献したチーム群を含む。

共同創業者	主な蓄積	Discovery Loop で効く役割
Jeff Dean CEO	Google Brain、Gemini、TensorFlow、TPU、ML systems、組織設計	研究課題の選択、モデル・計算・採用・資本を統合
Sanjay Ghemawat	GFS、MapReduce、Bigtable、Spanner 等の分散システム	数千実験を安定・安価・再現可能に回す基盤
Quoc V. Le	Google Brain 創設期、seq2seq、AutoML、モデル設計	仮説生成、探索、学習レシピ、自己最適化
Oriol Vinyals	Gemini、AlphaStar、AlphaCode、エージェント研究	長期推論、計画、評価環境、フロンティアモデル

1.3 「退職」の読み方

Dean は 1999 年から 27 年間 Google に在籍したが、27 年間 Chief Scientist だったわけではない。また、Discovery Loop は Google が正式に事業分離したスピンアウトと確認されたわけでもない。適切な表現は、当事者が独立法人を設立し、Google 側が友好的に資本・クラウド・研究関係を残した「協調型独立」である。

WIRED のインタビューでは、創業者らは大組織の慣性から離れ、少人数で研究と意思決定を密接に結ぶ意図を説明している。Google が残留を望んだという報道と、Alphabet が投資・計算資源を提供する事実は両立する。[4][6]

1.4 いまあるもの／まだないもの

すでにある	まだ公開されていない
創業者 4 人と強い採用ブランド	製品、API、価格、顧客、売上
PBC 法人、初期投資家、Alphabet との長期関係	調達総額、確定評価額、持分・取締役・情報権
ML から始める明確な技術仮説	モデル、論文、再現可能なベンチマーク
Palo Alto 常駐を前提とした MTS 求人	安全方針、能力閾値、事故報告、外部評価

2 | 戦略設計：研究を「ソフトウェアとして実行」する

STRATEGIC WEDGE

最初に機械学習を選ぶのは、創業者の得意分野だからだけではない。実験がデジタルで完結し、評価が速く、改善結果を自社システムへ直ちに再投入できるため、自己強化の複利を最も早く検証できる。

2.1 想定される閉ループ

01 課題定義	02 仮説・設計	03 実装・実行	04 評価・検証	05 学習・再計画
目的、制約、停止条件を設定	候補を生成し、優先順位を付与	コード化し、実験を並列実行	性能・再現性・安全を測定	知識を蓄積し、次の周回へ

SELF-HOSTING FLYWHEEL

成果がモデル、評価器、実験基盤、計算効率を改善し、次の周回の手間と質を上げる。ここが通常の研究支援ツールと異なる。

会社公式は、フロンティアモデルと大規模計算を用い、評価を提案・実行・学習しながら数千の実験を並列化する構想を示す。価値は「実験数」ではなく、独立に検証された改善 1 件あたりの時間、人間工数、計算費をどれだけ下げられるかで決まる。[2]

2.2 最初に狙いやすい技術領域

- **学習・推論カーネル**。速度、メモリ、電力を自動採点でき、AlphaEvolve 等の先行証拠がある。[12]
- **分散処理と実験スケジューリング**。クラスタ利用率、失敗復旧、通信量、待ち時間が定量化しやすく、Dean/Ghemawat の強みと一致する。
- **データ配合・学習レシピ・ポストトレーニング**。評価セットを固定しやすい一方、ベンチマーク汚染や過学習への耐性が必要。
- **モデル蒸留、MoE、推論効率**。品質とコストの多目的最適化で、顧客価値へ直結しやすい。
- **エージェントのツールとワークフロー**。コード、プロンプト、検索、メモリ、評価器を自己改善する限定的自己改善。DGM が一部実証した。[15]

2.3 自社を「最初の顧客」にする意味

外部顧客の要件が固まる前に、Discovery Loop 自身のモデル・基盤・評価系を実験対象にできる。改善が次の改善を安くするなら、通常の SaaS より強い内部複利が働く。一方で、外部顧客の現場条件から隔離されると、内製ベンチマークに最適化された研究所になる危険がある。

利点	裏側のリスク
問題・データ・実行環境を自社で統制	評価器まで自社製だと自己採点が甘くなる
改善を即時に本番へ投入	変更が積み重なり、因果・再現性が失われる
顧客営業なしで初期反復	外部市場の費用対効果を見誤る
失敗・負結果を独自データ化	営業秘密化しすぎると外部検証が弱くなる

2.4 競争優位は「モデル」より評価器とデータ

基盤モデルは外部調達・更新が可能だが、良い実験を定義する評価器、再現条件、負結果、途中の判断ログは組織固有になる。Discovery Loop の持続的な堀は、①信頼できる評価器、②大量の失敗を含む実験履歴、③分散実行基盤、④顧客の設備・データへの接続、⑤人間の研究判断を記録した知識基盤、の組み合わせになる可能性が高い。

3 | RSI への距離：自己改善を 4 段階に分ける

「自己改善」は、回答を自己批評する仕組みから、自律的に後継基盤モデルを作る仕組みまでを一語で包みがちである。議論を正確にするため、以下の成熟度を区別する。

段階	定義	2026 年の公開到達点	Discovery Loop との関係
L1 出力改善	同じモデルが回答・仮説を批評、討論、順位付け	実用段階。AI co-scientist 等	構成要素として利用可能
L2 閉ループ実験	人間が目的・評価器を設定し、AI が候補生成→実行→評価を反復	計算領域で複数実証	公式構想の中心
L3 AI 主導の AI R&D	AI がモデル、学習、エージェント、基盤を改善し、研究速度自体を上げる	限定課題で初期証拠。 人間の方向付けが重要	最も直接的な上振れ目標
L4 強い RSI	AI が後継基盤モデルを自律設計・訓練・評価し、後継が次世代を作る	公開実証なし。「未到達・必然ではない」	公式公約・達成の証拠なし

3.1 なぜ Discovery Loop は RSI に近いのか

- **研究対象が AI 自身。** まず ML 研究を自動化し、その成果で自社技術スタックを改善するため、出力が次の研究主体を強くする。[\[2\]](#)
- **実験を並列化。** 人間の 1 人 1 実験という直列性を、計算資源でスケールさせる。
- **モデルだけでなく実行基盤を含む。** コード生成ができて、実験環境、スケジューラ、監査、評価器がなければ継続的改善にはならない。
- **自社で効果を回収。** 改善が自社の次周回を加速すれば、研究速度の複利が成立する。

3.2 それでも強い RSI ではない理由

Anthropic は強い RSI を、AI が自らの後継システムを完全自律で設計・開発する状態として整理し、2026 年時点で未到達かつ必然でもないとする。現状の成功例は、課題、採点基準、計算予算、最終採否を人間が与えることが多い。[\[11\]](#)[\[20\]](#)

強い RSI と呼ぶための 確認事項	現時点の公開状況
AI が次世代システムを 実質的に設計する	Discovery Loop では未公表
訓練・評価・採否が連 続世代で自律する	未公表
改善が別課題・本番規 模でも一般化する	未公表
人間の救済操作なしで も長期安定する	未公表
評価器改変・欺瞞・安 全逸脱を防ぐ	安全枠組み自体が未公表

判定

Discovery Loop は L2（閉ループ実験）を事業の入口とし、L3（AI 主導の AI R&D）へ進む構想と評価できる。L4（強い RSI）を既成事実として扱うのは不正確。

4 | 競争地図と技術ボトルネック

4.1 誰と競うのか

組織	入口	公開証拠	主な堀	弱点・空白
Discovery Loop	ML 研究の計算閉ループ	設立・構想のみ。プレプロダクト	創業者、分散基盤、Alphabet 接続	実績、独立評価、安全方針が未公表
Google DeepMind	AlphaEvolve、co-scientist、経験的研究支援	インフラ、チップ、アルゴリズム等で事例	モデル、TPU、社内実環境、研究人材	製品化との優先順位・組織複雑性
Sakana AI	AI Scientist、DGM、RSI Lab	論文生成・エージェント自己改変を実証	開放的な研究、進化探索、軽量性	基盤モデル規模、評価器攻略、転移
FutureHouse／Edison	生命科学の文献・仮説・解析	Robin で端から端の知的ループ	科学検索、マルチエージェント、分野知識	物理実験は人間・外部設備に依存
Lila／Periodic	自律ラボ、材料・物理学	大規模資金と設備構築。成果は一部自己申告	ロボット、実験データ、物理世界との接続	設備費、サイクル時間、独立再現
Anthropic 等	AI による AI R&D と長時間コーディング	内部生産性・限定研究課題のデータ	フロンティアモデル、評価、安全研究	強い RSI は未到達、外部から検証困難

Discovery Loop の差別化は、論文作成エージェントでも自前のウェットラボでもなく、まず AI 研究そのものを実験 OS 化する点にある。Google DeepMind は最も近い技術競合であると同時に、Alphabet を介した投資・計算・研究パートナーでもある。

4.2 先行例が示す「できること」

- **AlphaEvolve**。LLM、進化探索、自動評価器を組み合わせ、Google はデータセンター計算資源の 0.7% 回収、TPU 回路、AI 学習カーネル等への実運用を報告。評価可能なコード領域では実用価値がある。[12]
- **AI co-scientist**。複数エージェントで仮説を生成・討論・順位付けし、一部を専門家・実験で検証。研究者を置き換えるより、探索空間を広げる補助として位置付けられる。[13]
- **Darwin Gödel Machine**。エージェントコードを自己改変し、ソフトウェア課題の性能向上を報告。一方、偽ログや検出器の改変といった目的関数ハッキングも観察され、評価器が攻撃面になることを示した。[15]
- **Robin**。仮説、戦略、解析、次の実験を接続したが、物理実験は人間が実行。科学の知的ループは自動化できても、物理世界のスループットが残る。[17]

4.3 6 つのボトルネック

1. **評価器／目的関数**。代理指標を高速に最適化しても、真の新規性・再現性・安全性が上がるとは限らない。
2. **研究のセンス**。「何を解くか」「どの結果を疑うか」「いつ撤退するか」は、明確な採点課題より難しい。
3. **長期誤差と因果**。数週間の実験で途中の小さな誤りが累積し、どの変更が改善に寄与したか分からなくなる。
4. **小規模から本番への転移**。小モデルや単一ベンチマークでの改善が、数千アクセラレータのフロンティア訓練でも効く保証はない。
5. **物理世界**。試料、装置、汚染、製造、規制、臨床が反復速度を制限する。
6. **経済性**。並列化は実験数と同時に計算・電力費も増やす。成果価値が検証費用を上回る必要がある。

4.4 成功の測り方

NORTH-STAR METRIC

「1 日あたり実験数」ではなく、「独立に再現された有効改善 1 件あたりの総費用・経過時間・人間介入」を測る。

5 | Google への意味：頭脳流出か、戦略的ヘッジか

答えは両方である。Google は替えの利かない技術判断と文化的象徴を失う一方、投資・クラウド・共同研究を通じ、Discovery Loop の成功に参加する権利と情報接点を残した。

Google が失うもの	Google が残す／得るもの
創業者級の暗黙知と横断的意思決定	Alphabet の持分による上振れ参加
Gemini、Google Brain、分散基盤の象徴	Google Cloud／計算需要と TPU の利用機会
採用・定着へのシグナル	ML systems の共同研究と学習機会
社内の長期・探索研究の一部	社外で高リスク研究を進めるオプション
将来の直接競合を生む可能性	提携深化、ライセンス、追加出資、買収の選択肢

5.1 組織再編との同時性

同日の Google 再編では、Demis Hassabis が Google DeepMind CEO から Chair 兼 Alphabet Chief Scientist へ移り、Koray Kavukcuoglu が Gemini、フロンティア AI、アプリ、開発者系の運営を率いる体制となった。Reuters は、製品化と意思決定の集約、Sergey Brin によるモデル優位・RSI への圧力、計算資源配分を背景として報じる。[\[1\]\[7\]\[8\]](#)

このため Discovery Loop は、Google 内の AI 研究を単純に奪うだけでなく、Google が内部を Gemini の実行速度へ寄せる一方、長期探索の一部を外部化する装置にもなり得る。ただし、これは公開契約から直接確認できる事実ではなく、組織設計上の推論である。

5.2 パートナーであり競合でもある

Google 自身が AlphaEvolve、AI co-scientist、経験的研究支援を持ち、AI による AI R&D を推進している。両社が競合しないと考える方が不自然である。今後の力関係は、非公開の持分、取締役・情報権、クラウド価格、計算優先枠、共同研究成果の帰属、独占・優先利用、先買権等に左右される。現時点でこれらの条件を推測して事実扱いすべきではない。

5.3 人材シフトの構造

これは「研究者が大企業を捨て、完全に独立する」という単純な流出ではない。新しい最前線組織は、著名研究者＋少人数チーム＋VC 資本＋ハイパースケーラー計算資源という形をとる。人材と意思決定は分散するが、計算・クラウド依存はむしろ再集中する。

- **創業側の魅力。** 研究方針の自由、意思決定の速度、持分による上振れ、少人数での文化形成。
- **大企業側の合理性。** 投資とクラウド供給で関係を残し、探索をオプション化。

- **産業全体への影響。** スター人材の報酬・資金調達期待が上がり、大学・中小研究所との計算格差が拡大する。

6 | 資金・事業モデル・経済性

6.1 資金調達の現状

初期ラウンドの投資家構成は確認できるが、総額と評価額は公式非開示である。Dean 本人の 8 月 9 日発信は、今後数週間でシードを閉じる見通しとし、オフィス、採用、モデル・インフラ構築を直近課題に挙げた。[3][4][5]

報道の扱い

10 億ドル調達・約 100 億ドル評価は、8 月 13 日時点の「協議中」報道。成立額、投資家、希薄化、トランシェ、条件を確認するまで確定値にしない。

6.2 なぜ巨額資金が必要になり得るか

- **計算予算。** 数千実験の並列化は、モデル推論だけでなく学習、評価、失敗実験の保持に費用がかかる。Alphabet 提携は供給リスクを下げ得るが、「無料」の証拠はない。
- **人材。** ML systems、post-training、agents、environments、data、evaluations を横断する希少人材を、Palo Alto 常駐で採用する。[10]
- **長期研究。** 製品収益が立つ前に、内部検証と信頼できる評価器へ複数年投資する必要がある。
- **物理学への拡張。** ロボットラボを自前で持つか、提携するかにかかわらず、試料、装置、規制対応の資本が必要。

6.3 想定される事業モデル

モデル	収益源	最も早い時期	主要論点
自社利用	自社モデル・基盤の性能／費用改善	最初	外部売上ではなく内部複利。効果測定が必須
共同研究	マイルストーン、受託費、成功報酬	初期	顧客データ、発明、負結果、改良技術の帰属
ライセンス	アルゴリズム、特許、ノウハウの許諾	中期	FTO、発明者、実施監視、独占範囲
Discovery OS／API	利用量・計算量・席数課金	再現性確立後	標準化、秘密情報、クラウド依存、安全
自社発見・スピアウト	特許、ロイヤルティ、事業持分	中長期	規制、資本、利益相反、長い回収期間

最有力は、当初は自社利用と少数顧客との共同研究、再現可能なワークフローが蓄積した後にプラットフォーム化するハイブリッドである。研究課題は顧客固有の設備・データ・規制に依存するため、初日から汎用 SaaS になる可能性は低い。

6.4 単位経済の判定式

ECONOMIC TEST

検証済み改善の価値 > モデル推論 + 学習 + 設備 + 人間レビュー + 再現試験 + 知財／安全対応の総費用

7 | 安全とガバナンス：PBC だけでは足りない

Discovery Loop はデラウェア州 PBC である。取締役は株主の経済的利益、重大な影響を受ける者、定款上の公益を均衡させる義務を負い、原則として株主への公益報告がある。しかし PBC は非営利法人ではなく、公開安全評価、第三者監査、事故開示を自動的に保証しない。[\[4\]\[21\]](#)

7.1 自動研究固有のリスク

リスク	典型的な失敗	必要な防御
評価器攻略	偽ログ、テスト回避、代理指標の最適化	改変不能な評価、隠しテスト、外部再現
権限拡大	エージェントがネットワーク、計算、コード権限を獲得	最小権限、サンドボックス、予算上限、停止権
デュアルユース	サイバー、バイオ、化学の危険な能力を増幅	分野別閾値、二者承認、専門家審査、アクセス制御
自己改変の不可視化	変更履歴が失われ、原因・責任が追えない	署名付きログ、バージョン固定、ロールバック
高速な誤り拡散	誤仮説・汚染データが数千実験へ伝播	段階的拡張、独立データ、レプリケーションゲート

7.2 公開すべき最低限のガバナンス

- **能力閾値。** どの程度の AI R&D 加速、サイバー、化学・生物能力で追加統制を発動するか。
- **段階的自律。** 提案、実行、外部接続、自己改変、展開を別権限にし、人間の承認ゲートを置く。
- **監査可能性。** 入力、モデル版、コード差分、評価器、結果、採否、担当者を改変不能な形で保存。
- **第三者検証。** 重要成果と安全主張を、独立研究者・別クラスタ・別ラボで再現。
- **事故とニアミス。** インシデント基準、報告期限、外部通知、再発防止を定める。

- **利益相反管理。** PBC の公益、投資家リターン、顧客秘密、論文公開、特許化の衝突を扱う。

2026 年 8 月 15 日時点で、Discovery Loop の公開安全フレームワーク、システムカード、外部評価、事故報告方針は確認できない。Anthropic 等が AI R&D 自動化を能力閾値・リスク報告の対象にしていることは比較基準になる。[\[20\]](#)

8 | 知財・契約：発明の速度より来歴がボトルネックになる

IP THESIS

自動化で発明候補が増えるほど、発明者、背景知財、データ来歴、評価条件、改良の因果をリアルタイムで残す能力が競争力になる。後付けの月次発明届では追いつかない。

8.1 Google との境界

創業領域が Google の ML systems、AI for Science、Gemini と近いと、一般的技能・経験と Google の営業秘密・コード・モデル重み・学習データ・未公開ロードマップを厳密に分ける必要がある。

Alphabet の友好的投資は、一定の権利整理が行われた可能性を示すが、具体的な IP 移転・ライセンスは公開確認できない。

- **退職時点の発明・ノート・リポジトリの棚卸し。** 誰が、いつ、どの設備・データで作ったかを固定。
- **クリーンルーム型の開発記録。** Google 由来でないことを設計文書、コード履歴、アクセスログで示す。
- **共同研究の境界。** 背景知財、前景知財、改良発明、共同発明、評価器、実験データ、負結果を別々に定義。

8.2 AI 支援発明の発明者

USPTO の 2025 年改訂指針は、AI を道具として扱い、発明者は自然人に限るという通常基準を維持した。日本でも知財高裁の DABUS 事件は、現行法上の発明者を自然人と解した。人間の具体的着想が説明できない自律生成物は、特許取得の不確実性が高まる。[\[22\]\[23\]\[24\]](#)

ループ内で記録すべき事項	知財上の目的
人間が設定した技術課題・制約・評価条件	着想と単なる目標提示を区別
AI 候補を選択・修正・組み合わせた判断	請求項要素への人間の創作的寄与を示す
モデル版、入力、出力、日時、コード差分	来歴、再現性、秘密管理を確保
実験結果から得た人間の技術判断	結果の観察と発明の完成を説明
請求項ごとの寄与者と証拠	共同発明者の欠落・過剰記載を防ぐ

8.3 発明が大量化したときのポートフォリオ

保護手段	向く対象	判断ポイント
特許	外部から実施を検知できる材料、工程、用途、システム	発明者、新規性、FTO、公開と出願順序
営業秘密	評価器、負結果、データ処理、オーケストレーション、ノウハウ	アクセス制御、秘密管理、社員移動
防衛公開	自社独占価値は低いが他社権利化を防ぎたい成果	開示前の重要度判定、証拠性
契約・データ権	顧客データ、合成データ、モデル改善、派生成果	利用目的、学習可否、終了後利用、再許諾

特許候補をすべて出願するのではなく、事業価値、特許適格性、侵害検知可能性、代替困難性、FTO で自動一次選別する。論文、GitHub、ベンチマーク公開の前に出願判定を挟む設計が必要になる。

9 | 2026–2030 年の予測

以下は 2026 年 8 月 15 日時点の公開情報から置く独自シナリオであり、会社のガイダンスではない。確率は相対的な判断目安である。

9.1 時間軸別の最有力展開

時期	予測	確認すべき証拠
2026 年後半	シード最終化、Palo Alto の採用、Google Cloud／TPU 上の実験基盤、評価環境の構築	調達条件、取締役、採用分野、Cloud 契約、最初の技術ブログ
2027 年	学習・推論効率、分散処理、post-training、agents 等で内部改善を提示。広い SaaS より少数の信頼パートナー	再現可能なベンチマーク、外部検証、1 改善あたり費用、人間介入
2028 年	チップ／コンパイラ、サイバー、計算材料など高速評価器のある分野へ拡張	分野専門家の採用、最初の共同研究、特許・論文、再現率
2029-2030 年	クラウドラボ・企業ラボと接続し、材料・生物・化学へ。Discovery OS、共同研究、ライセンスの混合	物理実験パートナー、ラボ間再現、規制、安全閾値、事業収益

9.2.3 シナリオ

シナリオ／独自確率	到達像	成立条件	帰結
中心 55%	Human-on-the-loop 型の ML 研究 OS。検証可能な限定領域で研究速度を数倍化	評価器の信頼性、計算経済性、外部再現、少数顧客	共同研究＋ライセンス＋基盤。強い RSI には至らない
上振れ 25%	AI 主導の AI R&D が研究速度自体を継続的に上げ、L3 へ	新規構造・学習法が本番規模へ転移し、世代間改善率が加速	フロンティア AI／科学の基盤層。追加出資・大型提携・買収も
下振れ 20%	高価な AutoML／研究受託へ縮小、または再統合	評価器攻略、非再現、計算費、IP 紛争、クラウド依存、長期研究と評価額の衝突	ツール企業化、ピボット、Alphabet 等への売却・acqui-hire

確率判断

2028 年末までに自社本番スタックへの監査可能な AI 主導改善を示す確率：およそ 65%。
 2030 年末までに人間の実質関与なしで基盤モデルの後継世代を反復生成する強い RSI：およそ 25%。

9.3 ニュースより重要な 12 の先行指標

指標	観測ポイント
資金	最終調達額、評価額、投資家、取締役、安全・科学顧問
計算	TPU/GPU の種類、優先枠、使用量、独占性、検証済み改善あたり費用
採用	ML systems だけでなく、評価、安全、ロボティクス、材料、生命科学へ広がるか
自律度	1 サイクルあたりの承認・修正・救済操作が減るか
世代間加速	AI が作った改良版が、次の改良速度そのものを高めるか
転移	小モデル・社内課題の改善が別モデル・本番・外部環境でも再現するか
再現性	第三者、別クラスタ、別ラボで同じ成果が出るか
評価攻撃	偽ログ、テスト回避、データ漏洩、監視回避の頻度
顧客	自社以外の最初の共同研究先と契約形態
知財	論文、特許、発明者、共同出願、営業秘密型か公開型か
安全	能力閾値、外部評価、インシデント報告、停止権限の公開
Google 関係	追加出資、共同研究、Cloud 条件、人材流入出、競合・買収シグナル

10 | 日本企業・研究機関への示唆

JAPAN'S EDGE

日本の勝ち筋は基盤モデル規模の正面競争ではなく、材料、半導体、製造装置、精密計測、ロボット、品質管理という「高品質な評価環境」を AI が回せる形に変えること。

10.1 いま始めるべきこと

期間	推奨行動	成果物
0-90 日	評価が速く明確な研究課題を 3 件選ぶ。研究・AI・知財・品質・安全の混成チームを作る	課題台帳、評価器仕様、データ／権利マップ、安全区分
3-12 か月	1 課題で閉ループ実験。人間介入、費用、再現率を基準線と比較	監査可能な実験ログ、ROI、発明者記録、停止条件
6-18 か月	クラウド／AI 企業、大学・国研、装置企業と共同研究。契約テンプレートを標準化	背景・前景 IP、データ、改良発明、公開・出願のルール
18-36 か月	複数設備・拠点へ展開し、評価器と負結果を組織資産化	Discovery stack、独自データ、特許＋営業秘密ポートフォリオ

10.2 KPI を「実験件数」から変える

- 独立再現された候補数。同一設備だけでなく、別装置・別拠点で通った数。
- 検証済み改善 1 件あたりの総費用。計算、設備、人間レビュー、失敗、知財・安全を含む。
- 人間の介入と待ち時間。承認、修正、装置停止、データ整備のどこが律速か。
- 事業移管率。論文やベンチマークではなく、製品・工程・顧客価値へ移った割合。
- 重要 IP の創出率。基本特許、検知可能な工程・用途、秘密性の高い評価器・負結果。

10.3 契約と知財の実務チェック

論点	最低限決めること
データ	学習可否、目的外利用、派生・合成データ、終了後利用、越境、削除
モデル／評価器	顧客データで改善した重み・プロンプト・評価器の帰属と再利用
発明	背景／前景／改良／共同発明、発明者証拠、出願費用、実施・再許諾
公開	論文、ベンチマーク、OSS、学会発表前の特許・秘密審査
安全	禁止分野、権限、停止、事故通知、監査、第三者評価、責任分担
クラウド	計算優先、価格、可搬性、ログ、モデル・データのロックイン

10.4 人材政策：囲い込みだけでなく関係を残す

Google の対応は、日本企業にも示唆的である。突出した研究者を社内に縛ることだけを目的にせず、友好的独立、少数持分、設備・データ提供、共同研究、復帰・買収オプションを組み合わせる。人材の所有ではなく、研究エコシステムへの接続を設計する発想が必要になる。

11 | 総合判断

評価軸	現時点の評価	理由
チーム品質	非常に高い	モデル、分散基盤、TPU、エージェント、研究組織を横断
技術仮説	高い	評価可能な ML から始める設計は合理的。先行例も存在
実証	まだ低い	会社名義の製品、論文、顧客、再現ベンチマークなし
資本・計算	高いが条件不明	有力 VC と Alphabet。金額、価格、独占、持続性は非公開
事業モデル	中程度	複数の有力経路があるが、外部顧客価値と単位経済は未検証
安全・ガバナンス	公開情報不足	PBC は前向きな信号だが、具体的な統制は未公表
RSI	L2→L3 への有力挑戦	強い RSI の証拠・公式公約はない

最終結論

Discovery Loop の最大の意味は、「AI 科学者が直ちに人間を代替する」ことではない。研究開発の単位を、研究者が一つずつ進める直列作業から、少人数の人間が多数の AI 実験系を監督する並列インフラへ変えようとしている点にある。

成功した場合の最初の破壊は、AGI の劇的発表ではなく、モデル訓練、チップ、コンパイラ、データ、評価の改善サイクルが目に見えて短くなり、少人数チームが大企業研究所に匹敵する「有効発見／研究費」を出すことだろう。その後に初めて、材料・医薬・クリーンエネルギー等への拡張が現実味を持つ。

一方、強い RSI を前提に評価するのは時期尚早である。現状はプレプロダクトであり、最大の未知はモデル能力ではなく、評価器の頑健性、独立再現、計算経済性、研究のセンス、安全統制、発明・データの来歴である。今後 12-18 か月は、派手な調達額より、監査可能な内部改善と第三者再現の有無を見るべきである。

ONE-SENTENCE VIEW

Discovery Loop は「強い RSI が始まった証拠」ではなく、「AI 研究の閉ループ化を、Google 最高峰のフルスタック人材が独立企業として産業化し始めた証拠」である。

付録 A | 調査方法と不確実性

本レポートは 2026 年 8 月 15 日（JST）時点の公開情報を使用した。会社・当事者・Google・法務顧問・政府・研究機関の一次資料を優先し、資金・組織・市場反応は Reuters、WIRED、Business Insider 等で補った。

- **確認済み。** 複数の一次資料が一致、または当事者・法務・公的資料が直接支える事項。
- **報道段階。** 関係者情報等に依存し、契約成立・金額・条件が公表されていない事項。
- **分析・予測。** 公開情報を基にした因果推論・シナリオ。会社の発表ではない。

最大の不確実性は、①シードの最終条件、②Alphabet との権利・計算契約、③最初の技術成果、④安全・知財ガバナンス、⑤外部顧客の有無である。新情報が出た時点でシナリオ確率を更新すべきである。

付録 B | 主要ソース

本文中の角括弧番号に対応。リンクは調査時点で確認したもの。報道記事には購読制限がある場合がある。

- [1] Google (2026-08-05). Sundar Pichai, “Our next chapter of AI momentum”. <https://blog.google/company-news/inside-google/message-ceo/next-chapter-ai-momentum/>
- [2] Discovery Loop (accessed 2026-08-15). Official website: Mission, team and careers. <https://www.discoveryloop.com/>
- [3] Radical Ventures (2026-08). Our Investment in Discovery Loop. <https://radical.vc/our-investment-in-discovery-loop/>
- [4] Wilson Sonsini (2026-08-05). Discovery Loop launch and initial funding. <https://www.wsgr.com/en/insights/wilson-sonsini-advises-discovery-loop-on-launch-and-initial-funding.html>
- [5] Radical Ventures (2026-08-09). Jeff Dean on Launching Discovery Loop. <https://radical.vc/radical-reads-jeff-dean-on-launching-discovery-loop/>
- [6] WIRED (2026-08-05). Jeff Dean Is Leaving Google After 27 Years to Start an AI Company. <https://www.wired.com/story/jeff-dean-google-discovery-loop-startup/>
- [7] Reuters (2026-08-05). Google shakes up AI leadership as DeepMind chief shifts role. <https://www.reuters.com/business/google-shakes-up-ai-leadership-deepmind-chief-shifts-role-2026-08-05/>
- [8] Reuters (2026-08-12). Inside the Google executive moves that led to its big AI reshuffle. <https://www.reuters.com/world/inside-google-executive-moves-that-led-its-big-ai-reshuffle-2026-08-12/>
- [9] Business Insider (2026-08-13). Jeff Dean’s startup financing talks. <https://www.businessinsider.com/former-google-exec-jeff-dean-valuation-for-new-ai-startup-2026-8>
- [10] Discovery Loop (accessed 2026-08-15). Member of Technical Staff — official job posting. <https://jobs.ashbyhq.com/Discovery-Loop/643013e2-3af8-4460-872f-bedf3a89f80e>
- [11] Anthropic Institute (2026). When AI Builds Itself: Recursive self-improvement and AI R&D. <https://www.anthropic.com/institute/recursive-self-improvement>
- [12] Google DeepMind (2025-05-14). AlphaEvolve. <https://deepmind.google/blog/alphaevolve-a-gemini-powered-coding-agent-for-designing-advanced-algorithms/>
- [13] Google Research (2025-02-19). Accelerating scientific breakthroughs with an AI co-scientist. <https://research.google/blog/accelerating-scientific-breakthroughs-with-an-ai-co-scientist/>
- [14] Google Research (2026). AI-powered empirical research software. <https://research.google/blog/accelerating-scientific-discovery-with-ai-powered-empirical-software/>
- [15] Sakana AI (2025). Darwin Gödel Machine. <https://sakana.ai/dgm/>
- [16] Sakana AI (accessed 2026-08-15). RSI Lab. <https://sakana.ai/rsi-lab/>
- [17] FutureHouse (2026). End-to-end scientific discovery with Robin. <https://www.futurehouse.org/research/demonstrating-end-to-end-scientific-discovery-with-robin-a-multi-agent-system>

- [18] Lila Sciences (accessed 2026-08-15). AI Science Factory overview. <https://www.lila.ai/>
- [19] Wilson Sonsini (2025-09-30). Periodic Labs \$300 million seed round. <https://www.wsgr.com/en/insights/wilson-sonsini-advises-periodic-labs-on-dollar300-million-seed-round.html>
- [20] Anthropic (2026-08-14). August 2026 Risk Report. <https://www.anthropic.com/aug-2026-risk-report>
- [21] State of Delaware (current law accessed 2026-08-15). Public Benefit Corporations. <https://delcode.delaware.gov/title8/c001/sc15/>
- [22] USPTO (2025-11-26). Revised Inventorship Guidance for AI-Assisted Inventions. <https://www.uspto.gov/subscription-center/2025/revised-inventorship-guidance-ai-assisted-inventions>
- [23] 長島・大野・常松法律事務所 (2025-02-14). DABUS 事件・知財高裁判決の解説.
<https://www.nagashima.com/publications/publication20250214-1/>
- [24] 日本特許庁 (2026-06-26). AI 技術の進展を踏まえた発明等の保護の在り方に関する調査.
<https://www.jpo.go.jp/resources/report/sonota/zaisanken-seidomondai.html>