

ANALYSIS REPORT | 2026 年 8 月

Discovery Loop の設立が示す AI 研究開発の転換点

ジェフ・ディーン氏らの独立、研究自動化、
再帰的自己改善（RSI）を多面的に読む

分析基準日 2026 年 8 月 15 日（GMT+9）

作成 Manus AI

対象 Discovery Loop / Google / AI for Science

本書は、公開情報に基づき、事実・報道・分析を明確に区別して整理したものです。

エグゼクティブ・サマリー

2026 年 8 月 5 日、ジェフ・ディーン氏はサンジェイ・ゲマワット氏、クオック・リー氏、オリオル・ビニャルス氏とともに、公益目的会社（Public Benefit Corporation, PBC）として **Discovery Loop** を設立すると公表した。同社が公式に掲げる使命は、機械学習、科学、工学の発見プロセスを自動化し、研究・開発の速度を高めることである。[1][2]

この動きを単なる「有名研究者の独立」とみなすのは不十分である。4 名の組合せは、大規模分散システム、計算インフラ、基盤モデル、ML 研究・製品化を横断する希少な**フルスタックの研究開発能力**を持つ。Discovery Loop は、まず ML 研究・エンジニアリングの実験ループを自動化し、自らを最初の顧客とする方針を示している。[2] これは、AI を科学の補助者にとどめず、AI 開発自体の試行回数・探索品質・実装速度を高める「研究生産システム」として設計する試みである。

もっとも、現時点で同社が明示しているのは**実験ループ自動化**であって、無制限の再帰的自己改善を公約しているわけではない。RSI に近い構図は、AI が ML 研究を改善し、その改善が次の AI 研究能力を増幅するフィードバックとして理解できるが、成立には自動評価の妥当性、再現性、計算資源、信頼できる実行環境、そして人間による統制が同時に必要である。[2][4][6]

本件の最も重要な含意は、巨大 IT 企業からスタートアップへ人材が一方向に流出することではなく、**研究者が新会社で探索速度と組織設計の自由度を得る一方、既存大手が資本・クラウド・基盤モデルを通じて関与し得る再編モデル**が現れている点である。Quartz は Google が創業投資家・クラウドパートナーになると報じたが、資本比率、契約期間、排他性は公表資料で確認できない。[5] したがって、Discovery Loop を「Google と完全に切れた競合」と断定するより、**部分的に補完しつつ将来は競合し得る、境界の流動的な研究開発ベンチャー**とみるのが妥当である。

結論	評価
直近の焦点	ML 研究・エンジニアリングの閉ループ自動化。科学全般への展開は中長期の目標である。[2]
最大の強み	分散システム、計算基盤、基盤モデル、プロダクトの結節点を理解する創業チーム。公式サイトは大規模計算と AI のフルスタック経験を強調する。[2]
最大の不確実性	資金規模、計算資源の契約、モデル戦略、最初の製品、外部顧客、ガバナンスは未公表である。

結論	評価
RSI の現状	「明示された完成目標」ではなく、ML 自動化を自社に適用する戦略から読み取れる 条件付きの近縁概念 である。[2][4]
今後 12~36 か月の勝敗	派手なデモよりも、独立再現可能な研究改善、計算効率、採用、外部検証、実行統制の 5 点で決まる。

1. 事実関係：確認済み情報と未開示事項

ディーン氏は自身の投稿で、ゲマワット氏、リー氏、ビニャルス氏と Discovery Loop を創業すると公表した。法人形態は PBC であり、使命は「機械学習、科学、工学を自動化して発見と進歩を加速する」こととされた。[1] 公式サイトは、提案、実行、結果の評価、次の反復という実験ループを、フロンティア AI モデルと大規模計算インフラにより自動化し、数千件の実験を並列に実行する構想を記している。[2]

重要なのは、「数千件を並列実行」は現時点の実測済み性能ではなく、会社が掲げるシステム設計・目標である点である。同様に、科学・工学へ広く展開する長期ビジョンはあるが、最初の対象は ML 研究・エンジニアリングである。同社はここで得た自動化能力を、まず自社技術スタックの最適化に用いるとしている。[2]

項目	公開情報の状況	本分析での扱い
共同創業者	ディーン氏、ゲマワット氏、リー氏、ビニャルス氏。公式発表・公式サイトで確認できる。[1][2]	確認済み
法人格	PBC として設立。[1][3]	確認済み
当初の技術対象	ML 研究・エンジニアリングの実験ループ自動化。[2]	確認済み
経営者	Radical Ventures はディーン氏を CEO と記載。TechCrunch も CEO 就任予定と報じた。[3][4]	高信頼の公表・報道
初期資金	Radical Ventures はシードラウンドの共同リード投資を公表。[3]	確認済み。ただし金額・評価額は未開示
Google との関係	Quartz は Google を創業投資家・クラウドパートナーと報道。[5]	二次報道。条件・比率・排他性は未確認
製品・顧客・収益化	公式サイトでは対面型の少数精鋭チームを構築中とする。[2]	未開示

項目	公開情報の状況	本分析での扱い
安全・ガバナンス	PBC であること以外に、公開の具体的な統制基準は確認できない。	未開示

解釈上の注意：PBC は、取締役会が株主価値だけでなく定款上の公益目的を考慮し得る法人設計である。しかし、PBC であること自体は、安全な研究運用、外部監査、または利益より公益を優先する個別意思決定を保証しない。本件では定款、取締役会、能力評価、外部助言、インシデント対応が未公表であるため、ガバナンスの実効性は判断できない。

2. なぜこの人材移動が重要か：4 人の結合が生む「研究生産システム」

Discovery Loop の差別化は、単一のモデル発明や論文テーマというより、**研究を実行する計算・ソフトウェア・評価・組織の全体最適化**にある。公式サイトは、創業チームが検索、分散データ基盤、TPU、TensorFlow、Gemini、AlphaFold 等に関わる経験を持つと説明し、チップからハードウェア基盤、ソフトウェア基盤、ML モデル、製品にまたがる深さを強調する。[2] 個別の功績の帰属は複雑であるが、少なくとも同社の組織的な強みが「大規模システムを研究成果から運用可能な製品へ移す能力」にあることは明確である。

従来の AI スタートアップは、基盤モデルの能力を API として利用し、特定業務への適用で価値を作ることが多かった。対して Discovery Loop が目指すのは、研究アイデア、実装、実験実行、評価、次の探索を一体化した**発見の工場**である。ML 研究を初期対象とすることで、湿式ラボや規制下の臨床試験に先行して、コード・学習ジョブ・ベンチマークという比較的自動化しやすい環境で、閉ループの性能を測ろうとしていると読める。

これは巨大 IT 企業にとっても二面的である。大手には基盤モデル、計算資源、データ、既存顧客、研究者群があるため、技術的に追従・代替しやすい。一方、大企業では既存製品との整合、収益責任、情報管理、組織の優先順位が研究の探索速度を制約する。独立組織は、少数精鋭で研究の目的関数と報酬を再設計し、失敗の許容度を高められる。Google の関与が報道どおりなら、Discovery Loop は「大手の資源」と「独立チームの探索速度」を組み合わせる、より複雑なモデルになる。[5]

観点	Discovery Loop にとっての機会	制約・反証可能性
人材	高い信頼関係を持つ共同創業者が、研究・インフラ・製品を横断して意思決定できる。	4 人の評判だけでは、エージェント・評価・データ・運用の組織能力は証明されない。
計算資源	大規模並列探索は、実験回数と探索幅を増やす。	学習・評価コストが成果の価値を上回る場合、規模は競争優位でなく負担になる。

観点	Discovery Loop にとっての機会	制約・反証可能性
自社適用	自社を最初の顧客にすることで、短い改善ループと実利用データが得られる。	内部ベンチマークの改善が、汎用的な研究能力や外部顧客価値を意味するとは限らない。
PBC	研究の長期性、社会的有用性、人材採用の物語を支え得る。	具体的な統治・検証・透明性がなければシグナルにとどまる。
Google との関係（報道）	資本・クラウドを通じた初期の実行力を高め得る。	依存度、知的財産、将来の競合・買収・人材採用を巡る緊張が生じ得る。

3. 技術の核心：実験ループ自動化と RSI を区別する

3.1 Discovery Loop が公式に述べるもの

Discovery Loop の公開説明は、研究目標から仮説・実験案を提案し、実装・実行し、結果を学び、次の実験を改善する**実験ループの自動化**である。[2] この構造は、AI が単発で文章やコードを生成することとは異なる。価値は、生成物を評価可能な実行に接続し、誤り・失敗・再現性を次の探索に取り込む点にある。

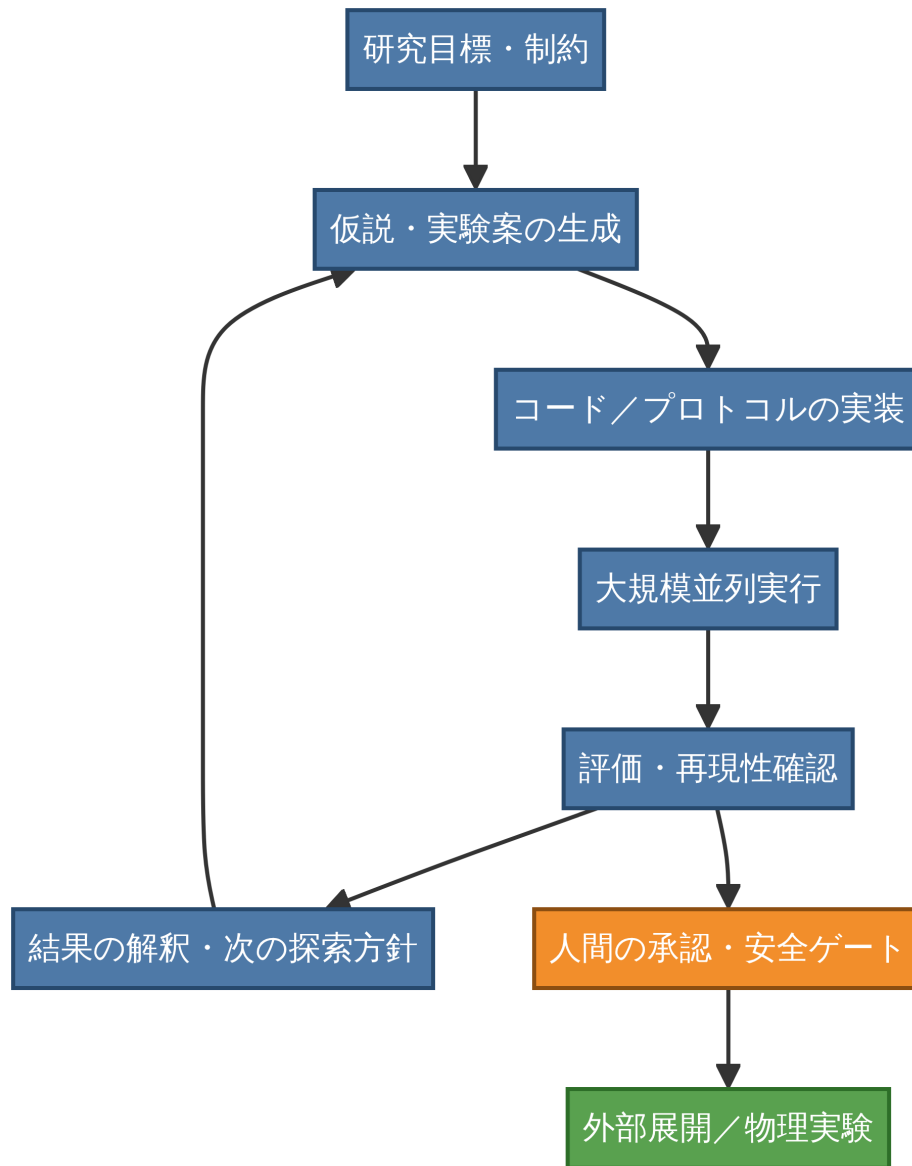


図1 Discovery Loop が掲げる構想を基にした、研究自動化の最小構成。図中の安全ゲートは、公式に公開された設計図ではなく、本分析が提案する運用上の統制点である。

ML 研究はこの仕組みの起点として合理的である。物理世界の実験に比べて、実行環境の複製、評価の自動化、ログ取得、反復の高速化がしやすい。2026 年の Nature 論文は、AI がアイデア生成、コード作成、実験、分析、論文執筆、自己査読まで行う研究自動化パイプラインを示した。ただし、生成論文はトップ会議の基準を一貫して満たしておらず、浅いアイデア、実装誤り、方法論的厳密性の不足、幻覚などが残ると報告している。[6]

3.2 RSI は「既成事実」ではなく、段階的な能力仮説である

RSI は一般に、AI が自らまたは次世代の AI システムを設計・改善することに意味のある貢献をし、その改善が次の改善能力を高める反復構造を指す。本件で重要なのは、Discovery Loop の公式発表は RSI という語を使わず、

まず ML 研究とエンジニアリングの自動化に集中するとしている点である。[1][2] TechCrunch は、同社が AI を用いてより強力な AI をつくることにも関心を持つと報じたが、これは公式の製品仕様・安全方針ではない。[4] 従って、現状を「自己改善 AI がすでに稼働している」と表現するのは誤りである。妥当な読み方は、**AI が研究プロセスを部分的に改善し、その成果が次の AI 研究を速める可能性がある設計**である。そこから強い RSI へ進むには、以下の条件を満たす必要がある。

成立条件	何を意味するか	失敗すると何が起きるか
評価可能な目的関数	ベンチマーク、計算効率、信頼性、コスト等を測定できる。	数値を最適化しても真の研究価値・汎化性能を損なう。
信頼できる実験実行	コード、データ、計算資源、実験環境が再現可能である。	失敗の原因がモデル・実装・環境のどれか分からず、ループがノイズ化する。
自動評価の妥当性	AI によるレビューや自己評価が人間の判断と整合する。	自己採点を「攻略」する報酬ハッキングが起きる。
改善の累積性	1 回の改善が、次の探索速度・探索品質を実際に高める。	局所最適、収穫逨減、計算費用の増大で回復が止まる。
統制可能な権限	実行できるコード、ネットワーク、計算、外部実験を段階的に制限する。	監査不能な実行、デュアルユース、逸脱した最適化が増幅される。

Google の AI co-scientist は、複数の専門エージェントによる仮説生成・反省・ランキング・進化・メタレビューを用い、専門家の関与下で生物医学の研究課題を扱った。[7] Sakana AI の AI Scientist も、計算実験で研究ライフサイクルの自動化を示す一方、論文生成の誤り、研究の浅さ、評価の限界を公開している。[9] これらは、Discovery Loop の方向性が技術的に突飛ではないことを示すが、同時に**自動化できることと、信頼して自律化できることの間には大きな隔たりがある**ことを示す。

4. 競争環境：競争相手はスタートアップだけではない

Discovery Loop の競争相手は、「AI for Science」を標榜する新興企業に限定されない。基盤モデルと計算資源を持つ巨大 IT 企業、研究自動化のオープンソース陣営、生命科学等の特定領域に深く入る組織、そして大学・国立研究所の自動実験プラットフォームが、異なる層で競争・協力する。

層	代表的な動き	強み	Discovery Loop への意味
---	--------	----	---------------------

層	代表的な動き	強み	Discovery Loop への意味
巨大 IT・基盤モデル企業	Google Research の AI co-scientist は、Gemini を用いるマルチエージェント型の研究支援を公表。[7]	モデル、計算資源、既存研究ネットワーク、配布力。	最も強い潜在競合であり、報道どおり Google との連携があれば当面は補完者にもなり得る。
AI 研究自動化プラットフォーム	Sakana AI の AI Scientist は、計算研究のアイデア生成から実験・論文化を公開し、コードも公開している。[9]	先行実装、コミュニティ、評価研究。	「AI が ML 研究を行う」こと自体は差別化にならない。大規模運用・成果の質・計算効率が争点となる。
科学領域に特化したエージェント	FutureHouse の Robin は、仮説・実験戦略・データ分析を統合し、生命科学の候補探索を報告した。[8]	ドメイン知識、実験協力、実データと検証。	物理世界へ展開する際には、領域別のデータ、実験設備、規制、検証体制が参入障壁になる。
物理実験・自律ラボ	化学、材料、バイオでのロボティクス・自動化設備。	実験を現実に行い、反復のデータを蓄積する。	Discovery Loop の長期ビジョンに不可欠な提携先・買収先・競合候補である。

この市場での勝ち筋は、最大モデルを訓練することでは必ずしもない。むしろ、(1) 実行可能な研究タスクを正しく分解し、(2) 実験を安く大量に回し、(3) 評価のゲーム化を防ぎ、(4) 人間の研究者が信頼できる再現可能な成果に変換する、という運用能力にある。特に科学・工学の外部展開では、データアクセスや実験装置の統合が、モデル能力以上の差別化要因になりやすい。

5. 事業化と組織設計：期待値を左右する 4 つの論点

5.1 「自社を最初の顧客にする」戦略の強みと罠

自社の ML 研究を対象にすると、外部顧客との連携や規制を待たずに、実験設計から実行ログ、評価までを一気通貫で測定できる。改善の効果を、学習コスト、性能、推論コスト、レイテンシ、信頼性、開発時間等で確認できるため、初期プロダクト・マーケット・フィットを社内で探索しやすい。[2]

他方、この戦略には循環論法の危険もある。内部で「AI が研究を改善した」と主張しても、評価器もデータも研究課題も自社が定義するなら、外部の科学的価値は保証されない。初期の信頼を得るには、事前登録された評価課題、外部ベースライン、第三者による再実験、失敗事例の開示が有効である。

5.2 資本・計算資源・独立性の三角関係

Radical Ventures はシードラウンドの共同リード投資を公式に公表した。[3] Quartz は Google が創業投資家・クラウドパートナーと報じた。[5] 報道が正しければ、Discovery Loop は初期の計算資源・設備・人材採用で有利になり得る。しかし、研究自動化の核心が計算・データ・開発ワークフローにある以上、クラウド依存、モデル提供者への依存、データの取り扱い、知的財産、競業避止の設計は事業上の重要な論点となる。

経営上の焦点は、Google との近さではなく、**どこまで独自のデータ・評価・実験基盤を蓄積できるか**である。大手のモデルを利用するだけなら、その大手が類似機能を製品化したときに交渉力は弱い。逆に、独自の研究運用データ、評価ハーネス、分散実験管理、専門チーム、再現可能な成果を積み上げれば、基盤モデルを横断して使える研究インフラ企業になり得る。

5.3 PBC は「意思決定の余地」であり、性能保証ではない

PBC という選択は、短期収益だけに縛られず、科学・工学の社会的利益を目的に含めるというシグナルになり得る。[1] [3] 人材採用、研究者・公的機関との協力、長期テーマへの資金配分にプラスに働く可能性がある。一方、重要な検証対象は、実際に公開される方針である。高リスク研究の事前審査、実行権限、計算資源のモニタリング、インシデント報告、外部専門家の関与、成果公開の基準が示されなければ、法人形態だけから安全性や公益性を推定することはできない。

5.4 「研究の自動化」は研究者を代替するのか

短期的には、研究者の仕事は消滅よりも再配分される公算が大きい。仮説の粗い探索、実験コードの作成、ハイパーパラメータ探索、文献の初期整理、結果の可視化は自動化率が上がる。その代わり、研究問題の定義、評価の妥当性、因果解釈、リスク判断、物理実験・臨床・制度との接続に人間の比重が移る。FutureHouse の事例でも、知的フレームワークを AI が担っても物理実験は人間研究者が行った。[8]

このため、最も起こりやすい組織変化は「研究者が不要になる」ことではなく、**少人数の高密度チームがより大きな探索面積を扱う**ことである。Discovery Loop がこのモデルを成功させれば、AI 企業だけでなく、製薬、材料、エネルギー、半導体、サイバーセキュリティの R&D 組織にも、研究人材・計算資源・実験設備の配置転換を促す可能性がある。

6. リスクとガバナンス：速さを価値に変える条件

研究自動化は、能力の加速と統制上の難しさを同時に増幅する。国際 AI 安全性報告書 2026 は、事前評価が実世界の有用性・リスクを確実に予測しない評価ギャップ、能力の予測不能な出現、開発速度とリスク管理の緊張を指摘する。また、脅威モデリング、能力評価、インシデント報告、多層防御を有効な実践として挙げる。[10]

Discovery Loop の文脈で特に重要なのは、AI エージェントが研究の「提案」だけでなく、「実装・実行・評価」へ近づくほど、権限管理が中核になる点である。計算実験であっても、コード実行、外部ライブラリ、ネットワーク、データ、クラウド予算、モデルの重み・API に対するアクセスを分離しなければならない。物理学・生命科学・サイバーセキュリティへ展開する段階では、実験設備、試薬、外部委託、危険情報、デュアルユースに関する別の安全ゲートが必要になる。

リスク	具体的な失敗様式	必要な統制・検証
評価のゲーム化	自己評価器や限定ベンチマークを最適化し、真の性能・新規性が伴わない。	外部ベースライン、隠しテスト、第三者再現、複数評価器。
誤りの大規模増幅	幻覚、誤実装、誤引用が並列実験と自動論文化で大量に増える。	実行前検査、来歴ログ、異常検知、抽出監査、公開時の AI 関与表示。
計算資源の濫用	失敗した探索が際限なく計算資源・予算を消費する。	予算上限、段階的実行、承認フロー、停止条件、単位成果当たりコストの開示。
デュアルユース	生命科学、化学、サイバー等の有害能力を高める。	領域別の能力評価、アクセス階層化、専門家審査、ログ監査、インシデント対応。
研究生態系への負荷	AI 生成論文・査読が増え、検証能力を上回る。	生成物の明示、データ・コード・実験ログの共有、査読制度との事前協議。

重要な判断：安全性は研究速度とトレードオフではなく、研究速度を持続可能な価値に変えるための生産品質である。再現不能な発見、監査不能な実行、外部検証不能な評価は、短期的なデモを作れても、科学・企業・社会の信頼を積み上げない。

7. 今後 12～36 か月：3 つのシナリオ

以下は公開情報と技術的条件を踏まえた**分析上のシナリオ重み**であり、確率を計測した予測ではない。資金規模、計算契約、採用状況、内部ベンチマークが未公表であるため、重みは新情報に応じて大きく変わり得る。

シナリオ	分析上の重み	2026 年末～2029 年前半に起こり得る姿	成立条件	早期の反証・警戒シグナル
------	--------	-------------------------	------	--------------

シナリオ	分析上の重み	2026 年末～2029 年 前半に起こり得る姿	成立条件	早期の反証・警戒シグナル
基準：ML 研究の高性能な内部基盤へ	50%	自社のモデル・インフラ・開発ワークフローを改善する実験基盤を構築し、一部を研究パートナーや企業へ提供する。	優秀な創業チームの採用、計算資源の確保、再現可能なベンチマーク改善、安定した実験運用。	技術発表がデモ中心で、対照実験・コスト・再現性が示されない。
上振れ：研究のフライホイルが実証される	20%	自動化が、自社の次世代モデル・システム設計を明確に改善し、独立した外部検証を伴う成果が複数出る。科学・工学の価値領域へ拡張する。	評価器の信頼性、強い実験オーケストレーション、独自データ・運用資産、物理実験パートナー、安全統制。	研究改善が単発で終わる、または大手の基盤モデル更新だけで成果が説明できる。
下振れ：創業者ブランドに運用が追いつかない	30%	採用・計算コスト・評価の難しさにより、汎用的な研究エージェントの差別化が弱まり、大手または既存プレイヤーに優位を取られる。	このシナリオを避けるには、顧客価値が明確な狭い課題を早期に選び、外部検証を積む必要がある。	継続的な研究成果・製品・人材採用が確認できず、計算資源依存だけが高まる。

最も蓋然性が高いのは、いきなり「科学を全面自律化する会社」になることではなく、ML 研究の一部を大幅に高速化する高性能な内部基盤を作り、その基盤を段階的に外部化する形である。大きな上振れは、AI が自社の研究・開発能力を継続的かつ検証可能に高めることを示せるかに依存する。これは RSI の完成を意味しないが、AI R&D の生産関数を変えるには十分に重要である。

8. 注視すべき指標：ニュースより実証を追う

今後の見極めで重要なのは、資金調達額や著名人の入社だけではない。以下の公開可能な証拠を追うことで、Discovery Loop が「注目企業」から「研究開発の新しい生産基盤」へ進んでいるかを判断しやすくなる。

指標	確認したい事実	なぜ重要か
技術的再現性	事前定義された課題、ベースライン、コード・ログ・コスト、第三者の追試。	実験数ではなく、発見・改善の質を判断できる。
自己適用の成果	自社の ML 開発で、性能、コスト、開発時間、信頼性をどう改善したか。	「自社を最初の顧客にする」戦略が実証されたかを測る。
外部検証	査読、独立研究者、顧客・共同研究先による評価。	内部評価のゲーム化や選択的開示を抑える。

指標	確認したい事実	なぜ重要か
計算効率	1 件の有用な改善・発見に要する GPU 時間、コスト、失敗率。	並列実行が経済的な競争優位かを判定する。
物理世界への接続	ラボ、実験設備、製薬・材料・エネルギー等との提携。	科学・工学への拡張可能性を検証する。
組織・ガバナンス	採用職種、PBC の目的・統治、能力評価、アクセス管理、外部助言。	野心的な自動化を安全かつ継続的に運用できるかを示す。
Google との関係	投資・クラウド・知財・モデル利用に関する正式発表。	独立性、資本効率、将来の補完・競合構造を左右する。

結論

Discovery Loop の設立は、AI 業界の注目人材が巨大 IT 企業を去ったという象徴的なニュースである以上に、**AI を「研究の成果物」から「研究そのものを実行する生産システム」へ変えようとする試み**として重要である。デューン氏らの強みは、AI モデルだけでなく、そのモデルを動かす計算資源、分散システム、評価、製品化までを一体として扱える点にある。[2]

ただし、研究自動化の実証はすでに競争的な分野である。Google の AI co-scientist、Sakana AI の AI Scientist、FutureHouse の Robin はいずれも、仮説、実験、評価の一部または全体を自動化する具体例を示している。[7] [8] [9] Discovery Loop の本当の勝負は、「AI エージェントを作れるか」ではない。**速く、安く、再現可能に、そして安全に、重要な研究課題を改善できるか**である。

RSI については、過大視も過小視も避けるべきである。公開情報から確実に言えるのは、同社が自らの ML 研究を自動化することで、AI 開発の実験ループを高速化しようとしていることだ。[2] これが持続的な能力増幅へ進むかは、評価の妥当性、実験の信頼性、計算経済性、外部検証、ガバナンスという、モデル性能だけでは解けない問題に依存する。ゆえに今後は、華やかな発表よりも、これらの条件を満たす公開の証拠を追うことが重要である。

分析上の前提・限界

本レポートは、2026 年 8 月 15 日（GMT+9）までに公開された一次情報、投資家発表、主要報道、査読論文、研究機関資料に基づく。Discovery Loop は設立直後であり、資金調達額、株主構成、クラウド契約、製品仕様、顧

客、独自モデル、計算資源、PBC の詳細な統治設計、安全方針は公開情報が限られる。Google が創業投資家・クラウドパートナーであるとの記述は Quartz の二次報道に基づき、条件の詳細は確認できない。[5]

本書のシナリオ重みは、公開情報と技術的要件に基づく分析上の見立てであり、投資推奨、企業価値評価、または個別の意思決定の助言ではない。将来の新たな一次発表、査読済み成果、資金調達、製品・顧客・ガバナンスの開示によって結論は更新されるべきである。

参考文献

[1] [Jeff Dean, Announcing Discovery Loop, 2026 年 8 月 5 日](https://x.com/JeffDean/status/2085034604172603724)

<https://x.com/JeffDean/status/2085034604172603724>

[2] [Discovery Loop — Continuous Exploration](https://www.discoveryloop.com/)

<https://www.discoveryloop.com/>

[3] [Radical Ventures, Our Investment in Discovery Loop, 2026 年 8 月 5 日](https://radical.vc/our-investment-in-discovery-loop/)

<https://radical.vc/our-investment-in-discovery-loop/>

[4] [TechCrunch, Jeff Dean and other top AI researchers are leaving Google to launch their own startup, 2026 年 8 月 5 日](https://techcrunch.com/2026/08/05/jeff-dean-and-other-top-ai-researchers-are-leaving-google-to-launch-their-own-startup/)

<https://techcrunch.com/2026/08/05/jeff-dean-and-other-top-ai-researchers-are-leaving-google-to-launch-their-own-startup/>

[5] [Quartz, Jeff Dean leaving Google after 27 years to co-found Discovery Loop, 2026 年 8 月 5 日](https://qz.com/jeff-dean-google-chief-scientist-discovery-loop-startup-080526)

<https://qz.com/jeff-dean-google-chief-scientist-discovery-loop-startup-080526>

[6] [Lu et al., Towards end-to-end automation of AI research, Nature, 2026](https://www.nature.com/articles/s41586-026-10265-5)

<https://www.nature.com/articles/s41586-026-10265-5>

[7] [Google Research, Accelerating scientific breakthroughs with an AI co-scientist, 2025 年 2 月 19 日](https://research.google/blog/accelerating-scientific-breakthroughs-with-an-ai-co-scientist/)

<https://research.google/blog/accelerating-scientific-breakthroughs-with-an-ai-co-scientist/>

[8] [FutureHouse, Demonstrating end-to-end scientific discovery with Robin, 2026 年 5 月 19 日](https://www.futurehouse.org/research/demonstrating-end-to-end-scientific-discovery-with-robin-a-multi-agent-system)

<https://www.futurehouse.org/research/demonstrating-end-to-end-scientific-discovery-with-robin-a-multi-agent-system>

[9] [Sakana AI, The AI Scientist: Towards Fully Automated AI Research, Now Published in Nature, 2026 年 3 月 26 日](https://sakana.ai/ai-scientist-nature/)

<https://sakana.ai/ai-scientist-nature/>

[10] [International AI Safety Report 2026](https://internationalaisafetyreport.org/publication/international-ai-safety-report-2026)

<https://internationalaisafetyreport.org/publication/international-ai-safety-report-2026>

— End of Report —