

Google DeepMindのAGIからASIへのロードマップに関する調査報告

Executive Summary

本件で日本語メディアが報じた「Google DeepMindがAGIからASIへの進化のロードマップを発表」という情報に最も対応する一次ソースは、Google DeepMindの公式公開ページとarXiv版の技術報告書“**From AGI to ASI**”である。arXiv版は2026年6月10日、DeepMindの公開ページは2026年6月12日付で、Tim Geneweinら14名が著者として記載されている。報告書の中核主張は、**AGIは終点ではなく通過点であり、AGIからASIへの移行には四つの主要経路がありうる**、という点にある。四経路は、**計算資源・モデル・データのスケールアップ、アルゴリズムのパラダイムシフト、再帰的自己改善、マルチエージェント集団の形成**である。さらに報告書は、これらの経路は**相互排他的ではなく、並行して進行しうる**と明示する。 ¹

ただし、DeepMindの文書は「予言」や「単一タイムラインの断定」をしているわけではない。むしろ逆に、**ASIの具体的な到達時期、実装詳細、確率評価は未指定**であり、報告書は「今の時点で、どれくらい速く能力が伸び、どこに天井があるかを信頼できる形で予測することはできない」と繰り返し強調する。時間に関して公式文書が示すのは、スケールアップ経路について「今後数年でさらに数桁の有効計算量やモデル規模の拡大は原理的には可能」とする程度であり、**AGI→ASIの暦年予測は示していない**。 ²

四経路のうち、DeepMind自身が「今日の視点では最も有望」としているのは**スケールアップ経路**だが、それは同時に最も古典的な「データ壁・資源制約・収穫逨減」の影響を受ける経路でもある。**パラダイムシフト経路**は必要性がありうる一方で内容が最も未指定であり、**再帰的自己改善経路**はAIによるAI研究自動化という強い仮説に依存し、**マルチエージェント経路**は分散的・組織的な超知能を視野に入れるが、協調コストや集団的誤作動の問題を抱える。学術文献を照合すると、各経路には相応の支持根拠がある一方、いずれも「単独で十分」とはまだ言えない。 ³

安全面で最も重要なのは、この報告書自体が「**AI Safety and Alignment will be solved to a sufficient degree**」という強い作業仮定の上に立っていることである。したがって、“**From AGI to ASI**”は**ポストAGIの能力進化地図ではあるが、それ自体が包括的な安全設計書ではない**。安全・制御・政策を論じるには、DeepMindの別文書—**Frontier Safety Framework**、**Technical AGI Safety & Security**、**AI Control Roadmap**、**multi-agent safety**研究助成—と併読する必要がある。これらの文書から導ける政策含意は、**単一のAGI到達判定ではなく、継続的な能力評価、事前安全審査、モデル流出防止、エージェント集団監視、国際協調**を組み合わせる必要がある、という点である。 ⁴

対象文書の特定と用語の整理

本調査で分析対象とした正式文書は、Google DeepMindの公開ページに掲載された“**From AGI to ASI**”と、その参照先である arXiv 論文 **arXiv:2606.12683** である。DeepMind公開ページは2026年6月12日付、arXiv登録は2026年6月10日で、著者は Tim Genewein, Matija Franklin, Alexander Lerchner, Laurent Orseau, Samuel Albanie, Adam Bales, Cole Wyeth, Stephanie Chan, Iason Gabriel, Joel Z. Leibo, Allan Dafoe, Marcus Hutter, Thore Graepel, Shane Legg と記載されている。報告書要旨は、AGI後のAI進展を「**機械知能の連続体**」の中で捉え、その主眼を **human-level AGI から artificial general superintelligence への移行**に置くと説明している。 ⁵

以下の日本語訳は、公式日本語版ではなく**筆者訳**である。まず題名は“**From AGI to ASI**” = 「**AGIからASIへ**」と訳すのが最も自然である。要旨の骨子は、「人間レベルのAGIは多くの大規模AI組織にとって次の10年の具体的な目標になっており、その実現後にAIがどのように発展しうるかを検討する。ASIを特徴づけたうえで、AGIからASIへの四つの経路と、その経路上の摩擦・ボトルネックを論じる。これらの摩擦が軽微か重大かは未解決で、多数の研究課題を生む」というものである。 ⁵

この報告書での**AGI**は、「単一の人間とほぼ同程度の知能をもつシステム」であり、より具体的には「多くの認知タスクで人間の中央値程度」と定義される。他方の**ASI**は、「人間の関心と活動のほぼすべての領域において超人的能力をもつ人工一般知能」であり、しかも単独の専門家ではなく、**大規模でよく協調した人間専門家集団を、ほぼあらゆる領域で上回る**という非常に高い基準が置かれている。また、ASIは単一個体とは限らず、数百万インスタンスの集合体であってもよいと明記されている。 ⁶

DeepMind報告の理論的な特徴は、AGIとASIの区別を厳密なベンチマーク閾値ではなく、**Legg-Hutter型の知能連続体**の上に置いている点である。さらにその上位概念として**Universal AI**を位置づけ、AIXIを理論上の上限として参照する。したがって、この報告書は「AGIとASIの間に断絶的な一回限りのイベントがある」と言うよりも、**知能の連続的上昇の中に複数の進化メカニズムが重なりうると**捉えている。 ⁷

日本語報道では、ビジネス+ITが2026年6月24日付で、DeepMind論文の主旨を「AGI到達を通過点として扱い、その後にASIへ至る具体的な四つの経路を提示した」と整理している。この報道は、四経路に加えて、AI進歩を継続的に測定・予測する仕組みの重要性も併せて紹介しており、原典の構造を比較的忠実に要約している。 ⁸

主要一次ソースのURLは次の通りである。

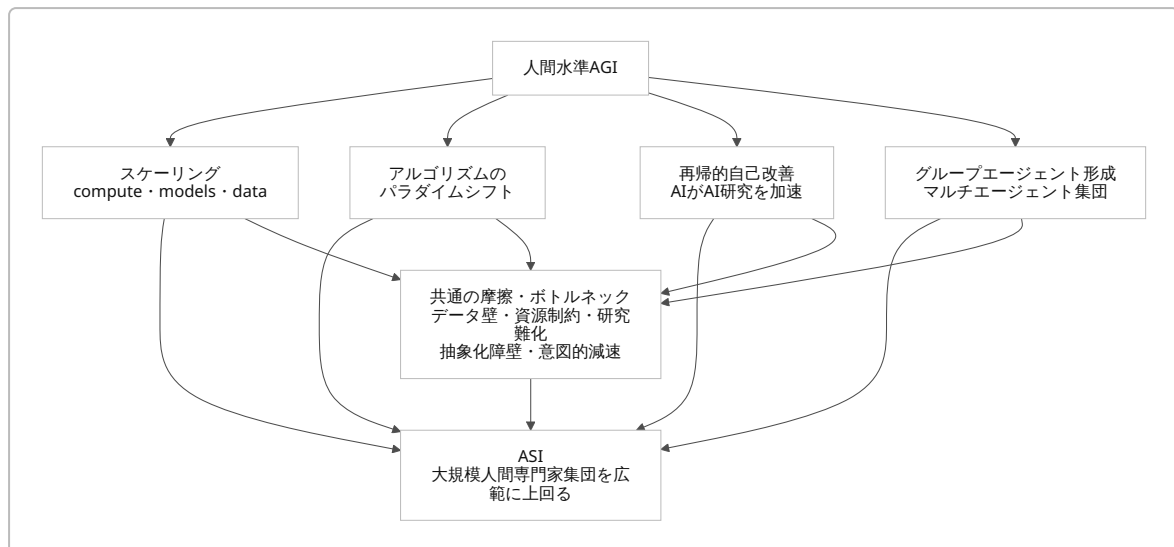
DeepMind公開ページ：
<https://deepmind.google/research/publications/239142/>

arXiv要旨：
<https://arxiv.org/abs/2606.12683>

arXiv PDF：
<https://arxiv.org/pdf/2606.12683>

四つの経路の正確な記述と比較

四経路の原文ラベルは、Table 3にそのまま整理されている。原文は“**Scaling compute, models & data**”, “**Algorithmic paradigm shift**”, “**Recursive (self-) improvement**”, “**ASI via group agent formation**”であり、DeepMindはそれぞれについて「主たる不確実性」も併記している。またTable 3の脚注では、これらの経路は**largely independent**で、しかも**並行して進行しうると**述べる。 ⁹



上図のとおり、DeepMindの図式は「一本道」ではない。たとえばスケーリングが限界に近づけばパラダイムシフトの圧力が上がり、再帰的自己改善が進めばスケーリング効率やアルゴリズム探索の速度が上がり、マルチエージェント化が進めば“個体”能力を超えた集団能力が出現しうる。つまり、この報告で言う“roadmap”は工程表というより、**相互作用する経路群の地図**である。 10

四経路の比較表

経路	正確な記述と筆者 訳	定義・技術メカ ニズム	前提	タイムラ イン	主な リス ク・ 摩擦	想定さ れる対 策・研 究課題	出典URL
ス ケー リン グ	“Scaling compute, models & data”／「計算資 源・モデル・デー タのスケーリン グ」 ⑨	既存の基盤モデ ル路線を延長 し、学習時・推 論時の計算資 源、モデル規 模、学習デー タを増やして性能 を押し上げる。 DeepMindは 「今後数年にわ たり有効計算量 とモデル規模を 数桁拡大できる 可能性」に言及 する。 ②	定量的ス ケー リン グが依 然とし て知能 向上に 効くこ と。ハー ドウェア・ ソフトウ ェア効率 改善が 継続す ること。 ②	未指定。 ただし 「今後数 年」にさ らに数桁 拡大の可 能性。 AGI→ASI の暦年予 測はなし。 ②	デー タ壁、 資源・ エネ ルギー 制約、 経済 性悪 化、収 穫逓 減、 抽象 化障 壁。 ⑪	テクノ 経済予 測、ポ スト人 間水準 ベンチ マーク、 合成デー タ・シ ミュ レー ション・ 自己生成 データ、 効率改 善。 ⑫	https:// arxiv.org/ pdf/ 2606.12683

経路	正確な記述と筆者 訳	定義・技術メカ ニズム	前提	タイムラ イン	主な リス ク・ 摩擦	想定さ れる対 策・研 究課題	出典URL
パラ ダイ ムシ フト	“Algorithmic paradigm shift”／「アルゴ リズムのパラダイ ムシフト」 ⁹	事前学習＋後学 習＋テスト時ス ケーリングとい う現行路線から 鋭く逸脱する新 しい学習原理・ アーキテク チャ・探索法で 進展する経路。 ¹³	現行 パラ ダイ ムが AGI/ ASI に不 十 分、 ある いは 限界 に達 する こと。 新パ ラダ イム が実 用的 利得 を示 すこ と。 ¹⁴	未指定。 DeepMind は「何が 新パラダ イムにな るかは予 測困難」 とする。 ¹⁵	シフ トは 十分 な規 模に なる まで 認識 され ない 可能 性、 研究 全体 の“低 い果 実”枯 渇、 技術 統合 コス ト。 ¹⁵	パラダ イム非 依存の 基礎理 解、理 論限界 と実用 限界の 把握、 新パラ ダイム を早期 検出す る評価 法。 ¹⁵	https:// arxiv.org/ pdf/ 2606.12683

経路	正確な記述と筆者 訳	定義・技術メカ ニズム	前提	タイムラ イン	主な リス ク・ 摩擦	想定さ れる対 策・研 究課題	出典URL
再帰 的自 己改 善	“Recursive (self-) improvement”／ 「再帰的自己改善」 9	AIがAI研究開発を高速化または自動化し、より良いAIがさらにAI研究を速める正のフィードバックを形成する経路。理論上は“explosive”な能力上昇もありうる。 16	AIが研究支援だけでなく、アルゴリズム改良、実験実行、データ生成などで実質的寄与を持つこと。 17	未指定。 急加速も早期頭打ちもありうる。とされる。 18	計算資源・実験・ハードウェア開発が物理的に時間を要すること、自己生成データ反復の劣化、収穫逓減。 19	再帰的改善のスケールリング則、AI R&D自動化の計測、ヒューマン・イン・ザ・ループ比率の追跡。 20	https://arxiv.org/pdf/2606.12683

経路	正確な記述と筆者 訳	定義・技術メカ ニズム	前提	タイムラ イン	主な リス ク・ 摩擦	想定さ れる対 策・研 究課題	出典URL
グ ルー プ エー ジェ ント 形成	“ASI via group agent formation”／ 「グループエー ジェント形成によ るASI」 ⁹	人間組織のよう に、多数のAGI エージェントの 協調から集団超 知能が生まれる 経路。中央集権 型から自己組織 化市場型までを 含む。 ²¹	集団 化に よる 知能 向上 が、 単体 モデ ルの 大型 化よ り効 率的 な問 題領 域が 存在 する こと。 ²²	未指定。 ただし DeepMind は近い将 来に「数 百万のAI エージェ ント」が 相互作用 する世界 を想定し ている。 ²³	計算 資 源・ エネ ル ギー 要 求、 協調 コス ト、 官僚 制的 オー バー ヘッ ド、 集団 的誤 作動 や予 測不 能な 新挙 動。 ²⁴	マルチ エー ジェン ト・ス ケーリ ング 則、サ ンド ボック ス、集 団監 視、評 判・ア イデン ティ ティ基 盤、人 口レベ ルの危 険検 知。 ²⁵	https:// arxiv.org/ pdf/ 2606.12683

各経路の技術的妥当性と批判的評価

スケーリング経路の評価

DeepMind自身は、現時点で最も現実的な経路としてスケーリングを位置づける。報告書は、近年のAI成功が「より大きなモデルを、より大きなデータで、より多くの計算資源で訓練したこと」によるとし、有効計算量とモデル規模を今後数年でさらに数桁拡大できる可能性に言及している。また、Suttonの“The Bitter Lesson”は、長期的には人手の知識注入よりも、検索と計算資源を活用する一般的方法が勝つと主張しており、DeepMindのスケーリング観と整合的である。²⁶

しかし、スケーリングだけでAGIあるいはASIに届くという見方には、一次資料ベースでも強い留保がある。DeepMind報告自体がTable 4で**Data wall**と**Economic and natural resource demand grows too fast**を主要摩擦に挙げている。外部研究でも、Villalobosらは、現行トレンドが続けば公開の人間生成テキストは2026年から2032年の間にほぼ使い切られると推計している。さらにShumailovらは、再帰的に生成データに再学習するとモデル崩壊が起きうると示し、AAAI 2025 Presidential Panelでは、回答者475人のうち76%が「現在のAI手法のスケーリングだけでAGIに至る」ことに懐疑的だった。²⁷

批判的に見ると、この経路の強みは「唯一、歴史データから予測モデルを当てやすいこと」だが、弱みは「その予測可能性が実際には資源・データ・抽象化能力の制約に強く依存すること」である。Nature系論考

でも、科学的発見の観点からは**スケーリング単独では十分でなく、解釈可能性・因果性・世界モデル・物理的検証との接続が必要**と指摘されている。したがって、スケーリングを主経路とみなすのは妥当だが、**それを単独経路とみなすのは過大評価**と言える。 28

パラダイムシフト経路の評価

この経路についてDeepMindは、現行の「大型基盤モデルの事前学習＋後学習＋テスト時スケーリング」から大きく逸脱する新しい原理が必要になりうると述べる。その際の不確実性はきわめて高く、何が有効な新パラダイムになるか、どんな計算・データ・エネルギー需要をもつかは**未指定**である。つまり、この経路は重要だが、ロードマップの中で最も“青写真が薄い”部分である。 29

もっとも、現行パラダイムの補完候補は周辺の公式研究から見えている。DeepMindは2026年3月に、AGIの進捗を認知能力の10分類で測るフレームワークを公開し、知覚・学習・記憶・推論・メタ認知・実行機能・社会的認知などの広い能力評価の必要性を打ち出した。さらに世界モデル研究については、Genie 3やD4RTを「AGIへの道の必要な一歩」と位置づけており、単なる次トークン予測を超える**環境表現・能動的相互作用・世界理解**を重視している。これは、将来のパラダイムシフトが「より強い世界モデル」「相互作用学習」「RL的学習」「物理接地」を含む可能性を示唆する。 30

批判的評価としては、**必要性は高いが、ロードマップとしての具体性は低い**。報告書は新アーキテクチャ名も、学習則も、移行条件も提示していない。そのためこの経路は、戦略論としては妥当でも、実務上は「研究の方向性」を示すにとどまり、予測可能性は低い。DeepMind自身が「基礎理解」と「理論限界の把握」を研究課題に挙げていることは、この経路がまだ概念段階に近いことの裏返しでもある。 15

再帰的自己改善経路の評価

再帰的自己改善は、AGIまたはそれ以前のAIがAI研究開発を加速し、その結果としてより強いAIがさらに研究を速めるという経路である。DeepMind報告は、このダイナミクスが「explosive」な能力上昇に至る可能性を認めつつも、歴史的先例がなく、爆発的成長から早期頭打ちまで幅広い可能性があるとする。これは、いわゆる「インテリジェンス・エクスプロージョン」を可能性としては認めるが、**実証的には未決着**とする立場である。 18

この経路を部分的に支持する実例として、DeepMindの**FunSearch**と**AlphaEvolve**は重要である。FunSearchは、LLMと自動評価器を組み合わせた進化的探索で数学・アルゴリズムの新発見を実現した。AlphaEvolveも、LLM群と評価器を用いてコードを反復改良し、科学・計算基盤の課題において既存手法を改善した。DeepMindのAGI安全保障文書も、AIが実験実装・データ作成・報酬モデリング等の研究サブタスクを加速し、AI研究そのものを速めうると明記している。これらは、**完全な自己改良ではないが、“AIがAI研究を助ける”という前段階はすでに実在すること**を示す。 31

他方で、批判的研究はかなり強い。WhitfillとWuは「計算資源ボトルネックがソフトウェアのみの知能爆発を妨げるか」を実証モデルで検討し、計算と労働が代替が補完かで結論が分かれることを示した。Chanらは、AI R&D自動化の進行度と効果を捉える既存ベンチマークが不十分だとし、資本支出、研究者時間、AIによる破壊的挙動の発生など、より制度的・実務的な指標を提案する。また、自己生成データの反復については、理論的にも退化ダイナミクスが生じうることが研究されている。総合すると、再帰的自己改善は**概念として非常に重要だが、いま現在の実証支持は「部分自動化」までであり、完全自律的RSIは未確認**である。 32

グループエージェント形成経路の評価

DeepMind報告の中で最も新味が強いのは、この経路である。単一モデルの知能向上だけでなく、**多数のAGIエージェントの協調が、組織・市場・企業のような集団主体を形成し、その総体としてASIに達する**という考

え方だ。報告書は、こうしたグループエージェントが、構成員とは異なる表象状態や動機状態をもつ可能性に言及し、中央制御型だけでなく自己組織化・分散市場型の出現も視野に入れる。これは従来の「単一超知能」像より、むしろ**制度・組織・経済系としての超知能**を強く意識した構図である。 33

この経路は、DeepMind周辺の安全研究とも整合的である。2025年の **Distributional AGI Safety** は、AGIが単一モノリスではなく、補完的なサブAGIエージェント群の協調から立ち現れる可能性を重視し、それに対応する分配的安全フレームワークを提案した。さらに2026年6月、DeepMindと複数機関はマルチエージェント安全研究に最大1000万ドル規模の助成公募を発表し、「数百万のAIエージェントが異組織間で通信・交渉・取引する」未来を想定している。つまり、グループエージェント経路は、同社の安全研究アジェンダの中でも**理論上の仮説ではなく、差し迫った設計課題**として扱われている。 34

一方で、この経路には明確な反証・反論もある。近年のマルチエージェント研究では、複数エージェント化が常に有利とは限らず、**coordination tax**、エラー伝播、文脈断片化、ツール利用との競合により、強い単一エージェントがマルチエージェント群を上回る条件も多いことが示されている。実際、ある整理では独立型マルチエージェントが単一エージェント比で17.2倍のエラー増幅を起こす一方、中央調整型でも4.4倍の残余があると報告する。また、集団通信そのものが偏極や“memetic drift”を生み、合意形成が「集合知」ではなく「集団的偶然」になるリスクも理論化されている。ゆえに、グループエージェント形成は非常にありうる経路だが、「多いほど賢い」という単純なスケーリング則は成立していない。 35

関連研究と反論の対照表

経路	支持する主要研究・事例	反論・制約を示す主要研究	評価
スケーリング	Sutton “Bitter Lesson”、DeepMindの「今後数年で数桁拡大可能」見解。 36	Villalobos のデータ制約推計、Shumailov のモデル崩壊、AAAI調査の懐疑、Nature論考の「scaling alone is insufficient」 。 37	短中期では主経路だが、単独では不十分な可能性が高い。
パラダイムシフト	AGI認知評価枠組み、世界モデル研究、相互作用学習への関心。 30	DeepMind自身が新パラダイムを未指定。予測も実装ロードマップも薄い。 15	必要性は高いが、具体性は最も低い。
再帰的自己改善	FunSearch、AlphaEvolve、AI R&D自動化の前段階。 38	Compute bottleneck論、AIRDA計測の未整備、自己生成データ退化。 32	部分的には進行中だが、完全RSIの実証は未到達。
グループエージェント形成	Distributional AGI Safety、DeepMindのmulti-agent safety助成、SPIRALの自己対戦学習。 39	coordination tax、error amplification、memetic drift、集団的誤作動。 35	非常に重要だが、制御と評価が最難関。

リスク、制御、安全対策、政策的含意

最初に強調すべき点は、“**From AGI to ASI**” 自体は安全計画を主題にしていないことである。報告書の研究課題一覧には、「AI Safety, Alignment, Sociocultural」として別枠が設けられ、そこでは**alignment が十分に解けていると仮定して本報告の考察を進めると**明言される。さらに、unsafe or uncontrollable なシステムは自動研究や大規模配備に活用しにくいため、整合性の難しさ自体が能力発展のボトルネックにもなりうるとしている。したがって、この文書だけを読んで「DeepMindはAGI→ASIの安全策をすべて提示した」と理解するのは誤りである。 40

DeepMindの実際の安全側文書を見ると、より具体的な制御策が提示されている。**Frontier Safety Framework** は、将来の重大危害をもたらす能力を **Critical Capability Levels** として特定し、早期警戒評価を定期的に回し、閾値到達時にはセキュリティ・デブロイ両面の緩和策を適用する枠組みである。2024年版はこの基本構造を示し、2025年版はセキュリティ水準の段階的対応や deceptive alignment リスクへの対応を強化し、2025年9月版では safety case review と Tracked Capability Levels を導入して、事前審査と継続監視をより制度化している。 41

AGI固有のリスクについては、DeepMindの **“An Approach to Technical AGI Safety and Security”** が、**misuse・misalignment・accidents・structural risks** を主要危険領域として整理し、とくに misuse と misalignment に重点を置く。misuse 向けには危険能力の評価、デブロイ・セキュリティ緩和策、レッドチーム的破壊試験を提示し、misalignment 向けには AI を使った監督、頑健学習、モニタリングとコンピュータセキュリティの defense in depth、alignment assurance を挙げる。この文書は、AGIに対する DeepMindの安全ロードマップが「存在する」ことを示すが、同時に未解決問題が残ることもはっきり認めている。 42

エージェント制御については、2026年6月の **AI Control Roadmap** が実務的に重要である。ここでは、内部AIエージェントを「潜在的に misaligned な insider threat」とみなし、**trusted AI supervisors による常時監視、harmful action の事前遮断、coverage・recall・time-to-response の測定、能力に応じた Detection/Response レベルの引き上げ**を掲げている。これは、再帰的自己改善経路とグループエージェント経路にまたがる **システムレベルの制御**として読むべきであり、報告書本文では未展開だった「どう抑え込むか」を補完している。 43

政策的含意は三つある。第一に、DeepMind報告が示すのは「AGI到達点」ではなく **継続的ポストAGI管理問題**であるため、規制も一回限りの認可ではなく、**反復的な評価・再評価・安全ケース更新**を前提にすべきである。第二に、グループエージェント経路が真剣に想定される以上、**単体モデル監査だけでは不十分**であり、人口レベルのエージェント相互作用、評判、トランザクション、群挙動の監視へ軸足を移す必要がある。第三に、報告書自身が指摘するように、ガバナンスは単なるブレーキではなく **進路制御**でもあるため、規制は capability cap 一辺倒ではなく、評価義務、モデル保護、事前配備審査、国際的リスク情報共有を組み合わせる方が実務的である。 44

現実の制度例としては、EUの一般目的AIモデル規制と Code of Practice は、GPAI提供者への義務をモデル単位で課し、systemic risk 判定の閾値や実務コードを整備している。国際面では **Bletchley Declaration** が、フロンティアAIのリスク認識とリスクベース政策構築の共有を打ち出した。米国では2026年6月の大統領令が、先端サイバー能力のベンチマークと“covered frontier model”の閾値判定を政府内で維持するよう指示している。DeepMindのロードマップを政策に落とすなら、こうした **能力閾値、事前評価、情報共有、国際協調**の組合せが最も整合的である。 45

結論と政策提言

結論として、Google DeepMindが2026年6月に公開した **“From AGI to ASI”** は、AGI到達後の能力進化を考えるうえで、かなり明確な枠組みを与えている。その核心は、**AGIからASIへの移行は一つの飛躍ではなく、四つの経路が並行・相互作用しながら進みうる**という認識である。四経路はいずれも技術的に一定の妥当性を持つが、一次ソースと学術文献を突き合わせる限り、**最も確からしいのはスケーリング、最も未指定なのはパラダイムシフト、最も劇的だが実証不足なのは再帰的自己改善、最も過小評価されがちだが制度的に重要なのはグループエージェント形成**である。 46

他方で、報告書を「ASI到来予言」や「具体工程表」と読むのは不正確である。**具体的実装、確率、到達年、単一路線、単純な知能爆発の必然性は、いずれも未指定**であり、DeepMind自身は不確実性の大きさを前面に置いている。さらに安全については、当該報告書が alignment 解決を仮定しているため、制御策を論じるに

は **Frontier Safety Framework**、**Technical AGI Safety & Security**、**AI Control Roadmap**、**multi-agent safety研究** を併せて読む必要がある。 47

政策提言としては、第一に、**AGIの有無ではなく post-AGI capability management** を制度設計の中心に置く必要がある。第二に、**ベンチマークと予測研究**をインフラ扱いし、人間水準を超えても有効な能力指標を整備するべきである。第三に、**モデル単体からエージェント集団へ**監督対象を拡張し、サンドボックス、評判・識別基盤、群挙動監視、人口レベルの危険兆候検出を進めるべきである。第四に、**能力閾値に応じた事前安全審査とセキュリティ水準引上げ**を法制度と企業内統制の両方で定着させるべきである。第五に、**国際協調**なしでは race dynamics によって単独規制が無効化されうするため、Bletchley 型のリスク合意と、EU・米国・研究機関の評価実務を接続する必要がある。 48

最後に、本報告における主な未指定事項を明記する。**ASI到達確率、AGI→ASIの具体年表、四経路の優先順位、単一実装アーキテクチャ、各経路の定量寄与、最終的な安全保証水準**は、原典では示されていない。したがって、最も厳密な読み方は、この文書を「未来予測の断言」ではなく、**ポストAGI世界に備えるための研究アジェンダとシナリオ地図**として扱うことである。 49

1 5 From AGI to ASI — Google DeepMind

<https://deepmind.google/research/publications/239142/>

2 4 9 10 11 12 13 14 15 16 17 18 19 20 24 25 26 27 29 40 47 48 49 From AGI to ASI

<https://arxiv.org/pdf/2606.12683>

3 6 7 21 22 28 33 44 46 From AGI to ASI

<https://arxiv.org/html/2606.12683v1>

8 Google DeepMindがAGIからASIへの進化のロードマップ発表 AGIからASIへ至る具体的な4つの経路を提示 | ビジネス+IT

<https://www.sbbbit.jp/article/cont1/185879>

23 Google DeepMind and partners announce multi-agent safety research funding call. — Google DeepMind

<https://deepmind.google/blog/investing-in-multi-agent-ai-safety-research/>

30 Measuring Progress Towards AGI: A Cognitive Framework

<https://deepmind.google/blog/measuring-progress-toward-agi-a-cognitive-framework/?lang=en>

31 38 FunSearch: Making new discoveries in mathematical ...

https://deepmind.google/blog/funsearch-making-new-discoveries-in-mathematical-sciences-using-large-language-models/?utm_source=chatgpt.com

32 Will Compute Bottlenecks Prevent an Intelligence Explosion?

https://arxiv.org/abs/2507.23181?utm_source=chatgpt.com

34 39 Distributional AGI Safety

https://arxiv.org/abs/2512.16856?utm_source=chatgpt.com

35 Towards a Science of Scaling Agent Systems

<https://arxiv.org/html/2512.08296v1>

36 The Bitter Lesson

https://www.incompleteideas.net/IncIdeas/BitterLesson.html?utm_source=chatgpt.com

37 [2211.04325] Will we run out of data? Limits of LLM scaling ...

https://arxiv.org/abs/2211.04325?utm_source=chatgpt.com

41 **Introducing the Frontier Safety Framework — Google DeepMind**

<https://deepmind.google/blog/introducing-the-frontier-safety-framework/>

42 **Taking a responsible path to AGI — Google DeepMind**

<https://deepmind.google/blog/taking-a-responsible-path-to-agi/>

43 **Securing the future of AI agents**

https://deepmind.google/blog/securing-the-future-of-ai-agents/?utm_source=chatgpt.com

45 **General-Purpose AI Models in the AI Act – Questions & Answers**

https://digital-strategy.ec.europa.eu/en/faqs/general-purpose-ai-models-ai-act-questions-answers?utm_source=chatgpt.com