

Educational Research Foundation
ERF Research Companion Series No.1

Google DeepMind

From AGI to ASI

第1回

AGIの先に何があるのか

なぜ Google DeepMind は「人間レベルの知能」で
はなく、その先を論じるのか

Publication Edition

Version 1.0

2026

目次

はじめに.....	2
第1節 この報告書の中心的な問い.....	3
第2節 AGIは「最終到達点」なのか.....	5
第3節 この報告書における AGI の意味.....	7
第4節 ASI とは何か.....	8
第5節 「Life as we don't know it?」という題名の意味.....	10
第6節 なぜ今、この問題を論じるのか.....	11
第7節 AGIは社会を根本的に変えるのか、それとも「通常の技術」なのか.....	12
第8節 計算資源の増大は、なぜ重要なのか.....	13
第9節 指数関数的成長と「シンギュラリティ」.....	14
第10節 この報告書が予言を避ける理由.....	16
第11節 この報告書の独特な「要約指示」.....	17
第12節 報告書全体の構成.....	19
第13節 この報告書の基本的な立場.....	20
第14節 高等教育にとっての意味.....	21
ERF Commentary.....	23
結論.....	25

はじめに

人工知能をめぐる議論では、長い間、AGI（Artificial General Intelligence：人工汎用知能）の実現が一つの到達点として語られてきました。

AGI とは一般に、特定の限定された課題だけではなく、人間が行う広範な知的活動を遂行できる人工知能を意味します。文章を書く、数学的問題を解く、科学的仮説を立てる、計画を作る、新しい状況に適応するといった、多様な認知的課題を横断して能力を発揮するシステムです。

しかし、Google DeepMind の技術報告書『From AGI to ASI』は、AGI の実現そのものを中心的な問題とはしていません。

この報告書が問いかけるのは、その一歩先です。

人間と同程度の汎用知能を持つ AI が実現した後、AI の能力向上はそこで止まるのだろうか。

あるいは、

AGI は最終地点ではなく、人工超知能 ASI へ向かう通過点にすぎないのではないか。

という問題です。

2026 年 6 月に公表されたこの報告書は、AI が人間レベルの AGI に到達した後も発展を続ける可能性を検討し、AGI から ASI へ至る技術的経路、その経路を加速または阻害する要因、そして現時点では答えの出ていない研究課題を体系的に整理しています。（Google DeepMind）

本稿ではまず、この報告書がなぜ書かれたのか、AGI と ASI をどのように位置付けているのか、そしてその問題設定が AI 研究だけでなく、社会、教育、科学、経済、政治にとってなぜ重要なのかを解説します。

第1節 この報告書の中心的な問い

『From AGI to ASI』の主題は、端的に言えば次の問いです。

人間レベルの人工知能が実現した後、人工知能はどのように発展していくのか。

著者たちは、AGIがいつ実現するかを正確に予測しようとはしていません。

むしろ、AGIの実現時期にかかわらず、いずれ人間レベルの汎用知能が実現したと仮定し、その後どのような技術的発展が起こり得るかを検討しています。

この報告書の目的は、将来を一つのシナリオに決めつけることではありません。AGI以後のAIの発展について、考え得る複数の経路を示し、それぞれの経路にどのような摩擦、制約、ボトルネックが存在するかを明らかにすることです。

報告書は、AGIからASIへの主要な経路として、次の四つを挙げています。

1. 計算資源、モデル、データの規模拡大
2. AIのアルゴリズムや学習パラダイムの転換
3. AIによる再帰的な自己改善
4. 多数のAIエージェントが形成する大規模な集合知

これらは互いに排他的な経路ではありません。複数の経路が同時に進み、互いを増幅する可能性があります。（Google DeepMind）

ここで重要なのは、著者たちが「ASIは必ず実現する」と主張しているわけではないことです。

また、「AGIが現れれば直ちにASIになる」と断定しているわけでもありません。

報告書の姿勢は、かなり慎重です。

ASIへの移行を加速させる要因と、それを遅らせたり停止させたりする要因の双方を検討し、現時点ではそのどちらが支配的になるか分からないとしています。

したがって、この報告書は ASI の予言書ではありません。

むしろ、AGI 以後の世界について、研究可能な問いを整理するための地図だと理解するのが適切です。

第2節 AGIは「最終到達点」なのか

これまでのAI研究や社会的議論では、AGIはしばしば究極の目標として扱われてきました。

人間と同程度に考え、学び、推論し、計画できるAIが実現すれば、AI研究は一つの完成地点に達するかのように考えられてきたのです。

しかし、この考え方には大きな疑問があります。

仮にAIが人間と同程度の知能を獲得したとして、その能力向上をそこで止める自然法則は存在するのでしょうか。

人間の知能水準は、生物進化によって偶然形成された一つの地点です。宇宙に存在し得る知能の理論的上限ではありません。

人間の脳の主な制約

人間の脳には、さまざまな物理的・生物学的制約があります。

- 神経信号の伝達速度
- 脳の大きさ
- エネルギー消費
- 記憶容量
- 学習に必要な時間
- 睡眠や疲労
- 寿命
- 同時に処理できる情報量

これに対してデジタル知能は、原理的には異なる性質を持っています。

同一のシステムを多数複製することができ、非常に高速に動作させることができ、膨大な記憶装置や外部ツールに接続できます。複数の AI を並列に稼働させ、知識や学習成果を共有することも可能です。

そのため、人間レベルの知能を実現することは、必ずしも AI 発展の終点ではありません。

それはむしろ、人工知能が人間の能力を参照基準とする段階から、人間の能力によっては測れない領域へ移行する転換点である可能性があります。

第3節 この報告書における AGI の意味

AGI には統一された定義がありません。

研究者や企業によって、AGI という言葉が示す範囲は異なります。人間と同等の知能を意味する場合もあれば、人間より優れた能力を持つ汎用システムまで含める場合もあります。

Google DeepMind は以前の研究『Levels of AGI』において、AI の能力を、能力の高さと一般性という二つの軸で分類する枠組みを提示しています。そこでは、性能、汎用性、自律性を区別し、AI の発展段階を比較可能にすることが試みられました。（Google DeepMind）

『From AGI to ASI』では、議論を明確にするために、AGI を比較的限定された意味で使用しています。

この報告書における AGI とは、おおむね、

広範な認知的課題において、平均的な一人の人間と同程度の知能を持つ人工システム

です。

ここで重要なのは、「最高の人間」ではなく「平均的または中央値の人間」が基準とされていることです。

また、AGI はあらゆる課題で人間と同じ性能を持つ必要はありません。

現在の AI も、すでにチェス、囲碁、タンパク質構造予測、計算、情報検索など、特定の領域では人間を上回っています。したがって、最初の AGI も、一部の課題では人間を大幅に超え、別の課題では人間に劣るという、能力の凹凸を持つ可能性があります。

報告書が重視するのは、個々の課題で均一な能力を持つことではなく、広範な認知的活動を横断して機能できる一般性です。

第4節 ASI とは何か

ASI は、Artificial Superintelligence（人工超知能）を意味します。

ただし、「ある一つの分野で人間より優れている AI」は ASI ではありません。

AlphaGo は囲碁で世界最高水準の人間を上回りました。AlphaFold はタンパク質構造予測において、人間の研究能力を大幅に拡張しました。しかし、これらは特定分野に特化したシステムです。

この報告書における ASI とは、

人間にとって重要なほぼすべての知的課題と領域において、人間を大幅に上回る汎用的な人工知能

を指します。

さらに、比較対象は一人の人間だけではありません。

報告書の要約では、ASI を、数千人規模の専門家集団が長期間協力して達成する能力をも上回るシステムとして説明しています。（arXiv）

つまり、ASI とは単なる「非常に頭のよい個人」の人工版ではありません。

それは、大学、研究機関、企業、政府機関のような大規模組織の知的能力を超える可能性を持つシステムです。

また、ASI が一つの巨大な AI である必要もありません。

数百万の AI インスタンスが並列に稼働し、それぞれが研究、計画、実験、評価、交渉、設計などを担当する集合体が、全体として ASI を形成する可能性もあります。

この点は極めて重要です。

一般に「超知能」という言葉からは、一つの巨大な人格的知性が想像されがちです。しかし、DeepMind の報告書が想定する ASI には、巨大な分散型組織や、多数のエージェントからなる集合的知能も含まれます。

第5節 「Life as we don't know it?」 という題名の意味

報告書の Introduction には、

Life as we don't know it?

という副題が付けられています。

通常、未知の生命を表現するときには、「life as we know it」、すなわち「私たちが知っている生命」という言い方が使われます。

それに対してこの表現は、私たちがこれまで経験したことのない知的存在や社会のあり方を示唆しています。

人類はこれまで、自分たちよりはるかに高度な一般知能と共存した経験を持っていません。

人間より高速に思考し、人間より多くの知識を保持し、疲労せず、複製可能で、互いに直接情報を共有できる知的主体が社会に存在したこともありません。

したがって、ASI の登場は、単なる新しい道具の登場ではない可能性があります。

それは、地球上の知的活動の構造そのものが変化する出来事になり得ます。

科学研究を誰が行うのか。

政策を誰が設計するのか。

企業を誰が経営するのか。

教育において、誰が知識を教え、誰が評価するのか。

文化や思想を誰が形成するのか。

こうした問いにおいて、人間が唯一の中心的な知的主体ではなくなる可能性があるからです。

第6節 なぜ今、この問題を論じるのか

著者たちは、過去 10 年間の AI 発展が極めて急速だったことを、この報告書の出発点としています。

AI 研究には、歴史的に見ても前例のない規模の計算資源、研究者、資金、企業組織が投入されています。フロンティアモデルは、以前は人間にしかできないと考えられていた認知的課題を次々と処理するようになっていきます。

報告書は、AGI が私たちの世代、場合によっては今後 10 年以内に実現する可能性を排除していません。（arXiv）

Google DeepMind 自身も、2025 年の安全性に関する文書において、AGI を「ほとんどの認知的課題で少なくとも人間と同程度に有能な AI」と位置付け、今後数年以内に実現する可能性に言及しています。（Google DeepMind）

ただし、ここで慎重が必要です。

「数年以内に実現する可能性がある」ということは、「数年以内に必ず実現する」という意味ではありません。

AI の進歩には大きな不確実性があります。

現在の技術路線がそのまま AGI へつながる可能性もあれば、データ、エネルギー、半導体、学習手法、推論能力、信頼性などの壁に直面する可能性もあります。

それでも、影響が極めて大きい以上、実現時期が不確実であることを理由に研究を先送りすることはできません。

AGI から ASI への移行に数十年かかるなら、人類には準備の時間があります。

しかし、AGI が AI 研究そのものを加速し、数年、数か月、あるいはさらに短期間で ASI へ接近するなら、AGI が実現してから対応を始めるのでは遅すぎる可能性があります。

この時間軸の不確実性こそが、この報告書を今書く必要性を生んでいます。

第7節 AGIは社会を根本的に変えるのか、それとも「通常の技術」なのか

この報告書は、AGIの社会的影響について一つの結論に固定していません。

一方には、AGIが経済、仕事、教育、科学、政治、文化、社会関係を根本的に変えるという見方があります。

AGIは一つの産業だけに適用される技術ではありません。ほぼすべての認知的活動に適用できる、汎用目的技術になり得ます。

蒸気機関、電気、コンピュータ、インターネットのように、多数の分野を同時に変えるだけでなく、それ以上に急速かつ広範な変化をもたらす可能性があります。

他方には、AIも最終的には「通常の技術」として社会に吸収されるという見方があります。

この立場では、AIはインターネットやスマートフォンと同様に大きな影響を与えるものの、社会制度、企業、政府、人々の行動が時間をかけて適応できる範囲にとどまります。

報告書はこの対立を決着させるのではなく、中心的な未解決問題として提示しています。

- AIの進歩は人間レベル付近で停滞するのか
- 人間集団を超えるASIへ進むのか
- 変化は数十年かけて進むのか
- 数か月から数年の急激な変化になるのか
- 社会は適応する時間を確保できるのか

この問いに答えるためには、AIの能力だけでなく、計算資源、エネルギー、アルゴリズム、組織、経済、規制、社会制度を同時に考える必要があります。

第8節 計算資源の増大は、なぜ重要なのか

報告書の導入部では、AIの進歩を支える三つの成長要因が説明されています。

第一は、半導体などのハードウェア性能の向上です。

第二は、計算設備への投資の増大です。

第三は、アルゴリズム効率の向上です。

アルゴリズム効率とは、同じ性能を達成するために必要な計算量が減少することを意味します。

同じハードウェアを使っているにもかかわらず、より優れた学習手法、モデル構造、最適化手法が開発されれば、実質的に使用可能な計算能力が増加したのと同じ効果が生じます。

報告書は、これら三つを組み合わせた「実効計算量」が、過去の傾向に基づけば非常に速い速度で増加してきたと論じています。ただし、その推計には大きな不確実性があり、将来も同じ成長率が継続する保証はありません。（arXiv）

重要なのは、個々のAIモデルの性能向上が止まったとしても、計算資源の増加によって、より多くのAIを同時に稼働させられる点です。

たとえば、人間と同程度のAIを千体稼働できる状態から、数年後に一億体稼働できるようになった場合、一体一体のAIが賢くなっていなくても、全体としての問題解決能力は大きく変わる可能性があります。

さらに、AIを人間より高速に動作させられるなら、研究や設計に要する時間も圧縮されます。

このため、AGIからASIへの移行は、一つのAIが突然天才になることで起こるとは限りません。

人間レベルのAIを大量に複製し、高速で並列稼働させるだけでも、集合体としては人間の大規模組織を超える可能性があります。

第9節 指数関数的成長と「シンギュラリティ」

報告書は、AIを支える実効計算量が一定の倍率で増え続ける場合、その成長は指数関数的になると説明します。

指数関数的成長では、増加量ではなく倍率が一定です。

毎年10増えるのではなく、毎年2倍、5倍、10倍になるような成長です。

指数関数的成長が続けば、初期には緩やかに見えても、後半には極めて大きな変化が生じます。

さらに、AIがAI研究そのものを加速するようになると、成長速度自体が上昇する可能性があります。

たとえば、

5. AIがより優れたAIを設計する
6. 改良されたAIがさらに効率よくAI研究を行う
7. そのAIが、さらに優れた次世代AIを設計する

という循環です。

これが再帰的自己改善です。

この循環が強く働けば、単なる指数関数的成長を超え、成長率そのものが上昇する超指数関数的な動態が生じる可能性があります。

これが、いわゆる知能爆発や技術的シンギュラリティと結び付けられる考え方です。

報告書は、I・J・グッド、レイ・ソロモノフ、レイ・カーツワイル、ニック・ボストロムなどの議論を参照しつつ、AIによるAI研究の自動化が急激な能力向上を引き起こす可能性を検討しています。(arXiv)

ただし、報告書はシンギュラリティを当然視していません。

現実の物理システムには、エネルギー、資源、製造能力、実験時間、データ、通信速度などの限界があります。

急速な成長も、通常はどこかで摩擦に直面し、成長率が低下します。

したがって、重要な研究課題は、「知能爆発が起こるか否か」を抽象的に論争することではありません。

どの段階で、どのような制約が現れ、どの程度成長を遅らせるのかを、具体的に検証することです。

第 10 節 この報告書が予言を避ける理由

AI の未来予測が難しい最大の理由は、加速要因と減速要因が同時に存在することです。

加速要因には、計算資源の増加、アルゴリズム改善、AI 研究の自動化、AI 同士の協働があります。

減速要因には、計算コスト、エネルギー供給、データ不足、半導体製造、研究課題の難化、物理実験に必要な時間、法規制、社会的抵抗などがあります。

将来の進歩は、この両者の競争によって決まります。

AI が研究を効率化する速度が、研究課題が難しくなる速度を上回れば、進歩は加速します。

反対に、新しい発見に必要な資源や実験時間の増加が、AI による効率化を上回れば、進歩は減速します。

どちらが勝るかは、現時点では分かりません。

しかも指数関数的または超指数関数的な成長を扱う場合、初期条件の小さな違いが、数年後には極端に大きな差になります。

そのため、単一の平均的な予測を提示しても、政策決定にはあまり役立たない可能性があります。

この報告書が目指すのは、未来の一点予測ではなく、複数のシナリオ、指標、ボトルネック、研究課題を整理することです。

第 11 節 この報告書の独特な「要約指示」

『From AGI to ASI』には、通常の学術論文には見られない特徴があります。

Section 1 の題名は、Summary Instructions——要約のための指示です。

著者たちは、人間の読者に対して、この報告書を AI アシスタントに要約させ、自分の関心や背景に合わせた説明を求めるよう勧めています。

さらに、AI アシスタントに対しても、どの論点を省略せずに説明すべきかを直接指示しています。

たとえば、

- AGI と ASI の定義を明示すること
- デジタル知能の利点を省略しないこと
- 四つの経路が同時進行し得ることを説明すること
- 摩擦やボトルネックを安易に圧縮しないこと
- 報告書公表後の進展を踏まえて評価を更新すること
- 報告書の欠点や見落としも示すこと

などが求められています。（arXiv）

これは単なる遊び心ではありません。

この構成自体が、知識の伝達方法が変わりつつあることを象徴しています。

従来の論文は、人間が固定された文章を読み、理解することを前提としていました。

しかしこの報告書は、AI が読者と報告書の間に入り、読者の専門性、関心、理解度に応じて内容を再構成することを前提としています。

報告書は固定された最終成果物ではなく、AI を通じて個別化され、更新され、批判される知識資源になりつつあります。

これは大学教育や学術出版にとっても重要な変化です。

第 12 節 報告書全体の構成

『From AGI to ASI』は、次のように構成されています。

構成	内容
Section 1	要約のための指示
Section 2	導入——AGI 後の世界をなぜ考える必要があるのか
Section 3	人工超知能 ASI の特徴
Section 4	Universal AI——理論上の知能の上限を考える枠組み
Section 5	AGI から ASI へ至る四つの技術的経路と、そのボトルネック
Section 6	補足的な論点と考察
Section 7	今後必要となる研究課題
Appendix A	報告書全体の要約
Appendix B	用語集

報告書の中心は Section 3 から Section 5 です。

まず ASI とは何かを定義し、次に機械知能の理論的上限を検討し、そのうえで AGI から ASI への具体的経路を論じます。(arXiv)

第13節 この報告書の基本的な立場

この報告書の立場は、楽観論にも悲観論にも単純化できません。

ASIが科学、医療、教育、経済、環境問題の解決を大きく前進させる可能性はあります。

一方で、人間の組織を超える能力を持つシステムは、制御、権力集中、安全性、社会秩序、雇用、意思決定の正統性といった重大な問題を生みます。

さらに重要なのは、社会が経験する変化が「AGIの登場」という一回の出来事ではない可能性です。

報告書は、AGIが登場した瞬間に一度だけ巨大な変化が起こるというイメージよりも、AIによる科学的・技術的突破が連続し、社会が繰り返し変革される可能性を示しています。（Google DeepMind）

すなわち、AGI革命が一回起きて終わるものではありません。

AIが研究、産業、教育、医療、エネルギー、ロボット工学を次々と変化させ、そのたびに社会制度の再編を迫る可能性があります。

第14節 高等教育にとっての意味

この報告書は、主としてAIの技術的发展を扱った報告書です。

高等教育そのものを中心テーマにはしていません。

しかし、その問題設定は大学の存在意義に直結しています。

これまで大学は、知識を保存し、生成し、伝達する主要な制度でした。

研究者は長い訓練を受け、専門分野の知識を蓄積し、論文や教育を通じて社会に還元してきました。

しかし、AGIやASIが実現すれば、次の前提が変化します。

- 知識を持つ主体は人間だけではない
- 研究を行う主体も人間だけではない
- 専門家集団より高度な分析をAIが行う可能性がある
- 個人に合わせた高度な教育をAIが提供できる
- 学問分野間の境界をAIが容易に横断できる
- 数百人、数千人規模のAI研究者を同時に稼働できる

この状況では、大学の役割を単なる知識伝達に置くことは難しくなります。

大学には、知識を教えること以上に、

- 何を研究すべきかを判断すること
- 知識の意味と社会的影響を吟味すること
- 異なる価値観の対立を扱うこと
- AIを含む多様な知的主体を統治すること
- 人間の尊厳、自律性、公共性を守ること
- 科学技術を社会的文脈に位置付けること

が求められるようになるでしょう。

AI が答えを生成できる時代には、大学の中心的使命は、答えを保有することから、問いを形成し、価値を判断し、知識を公共的に統御することへ移行する可能性があります。

ERF Commentary

Google DeepMind の『From AGI to ASI』が持つ最大の意義は、AGI を一つの完成地点として考える思考を転換したことにあります。

これまで私たちは、「AGI はいつ実現するのか」という問いに注目してきました。

しかし、本当に重要なのは、

AGI が実現した後、知能、社会、教育、科学がどのように再編されるのか

という問いです。

AGI が平均的な一人の人間に相当する知能であるなら、その複製と連携だけでも、大学、研究機関、企業、行政機関に匹敵する人工的組織が成立する可能性があります。

さらに、その組織が AI そのものの研究を行えば、知的進歩の速度も変化します。

このとき大学は、AI に代替されるか否かという受動的な問いだけを立てるべきではありません。

むしろ、

- 大学は ASI 時代の知識秩序をどのように形成するのか
- 人間と AI による研究共同体をどう設計するのか
- AI が生成した知識の正統性を誰が判断するのか
- 技術的に可能なことと、社会的に望ましいことをどう区別するのか
- 人間の成長、人格形成、公共的理性をどのように維持するのか

という問いを先取りする必要があります。

Ronald Barnett が論じる Ecological University は、大学を閉じた知識生産組織ではなく、社会、自然、制度、知識、文化などの複数の生態系に責任を持つ存在として捉えます。

ASI 時代の大学にも、これと同様の生態学的視点が必要です。

AIを単なる教育ツールとして導入するだけでは不十分です。

AIが科学、経済、政治、文化、人間関係を同時に変えていく中で、大学はそれらの相互関係を理解し、異なる価値と利害を調整し、技術の方向性を公共的に問い直す場でなければなりません。

その意味で、『From AGI to ASI』はAI技術の未来を論じた報告書であると同時に、大学が将来どのような知的・倫理的・公共的役割を担うべきかを考えるための重要な出発点でもあります。

結論

Google DeepMind の『From AGI to ASI』は、人間レベルの人工知能を AI の発展の最終地点とは見なしません。

AGI は、より広大な機械知能の連続体における一つの地点です。

その先には、計算規模の拡大、アルゴリズム革新、自己改善、多数の AI による集合知を通じて、人間の個人だけでなく、大規模な人間組織の知的能力をも上回る ASI が成立する可能性があります。

しかし、その経路も速度も確定していません。

AI の発展を加速する力と、それを抑制する物理的・経済的・技術的・社会的制約が、同時に作用するからです。

だからこそ、必要なのは断定的な未来予測ではありません。

複数の可能性を検討し、進歩を測る指標を整備し、ボトルネックを特定し、技術、安全性、社会制度、倫理、教育を横断する研究を進めることです。

この報告書の根底には、アラン・チューリングの言葉が置かれています。

私たちは遠くまで見通すことはできない。しかし、目の前には、なすべきことが数多く見えている。

『From AGI to ASI』は、ASI がいつ到来するのかを予言する報告書ではありません。

それは、見通すことのできない未来を前にしながらも、今から研究し、準備し、議論しなければならない課題を示す報告書なのです。