

Educational Research Foundation
ERF Research Companion Series No.1

Google DeepMind

From AGI to ASI

第8回

AGI から ASI への進歩を 妨げるもの

データ、資源、研究の難化、抽象化の壁、
身体性、そして社会による減速

Publication Edition

Version 1.0

2026

目次

はじめに.....	2
第1節 六つの主要なボトルネック.....	4
第2節 第一のボトルネック——データの壁.....	5
第3節 第二のボトルネック——経済的・自然資源需要の急増.....	9
第4節 第三のボトルネック——現在のニューラル・パラダイムの限界.....	13
第5節 第四のボトルネック——研究の難化.....	16
第6節 第五のボトルネック——抽象化の壁.....	19
第7節 身体性・現実接地のボトルネック.....	22
第8節 第六のボトルネック——意図的な減速.....	25
第9節 六つの制約はどのように重なり合うのか.....	28
第10節 AGIでちょうど進歩が止まる可能性.....	30
ERF Commentary.....	32
第11節 大学はAIの摩擦を減らすだけの制度ではない.....	32
第12節 データの壁と大学の知識資源.....	32
第13節 抽象化の壁と人間の研究.....	33
第14節 身体性と大学.....	33
第15節 研究の速度と学問の時間.....	34
第16節 意図的な減速と大学の公共的責任.....	34
第17節 大学は「アクセル」と「ブレーキ」の両方を担う.....	35
第18節 Ecological University との関係.....	36
第19節 ボトルネックを取り除くことが常に善とは限らない.....	36
結論.....	37

はじめに

これまで、Google DeepMind の『From AGI to ASI』が示す四つの技術的経路を検討してきました。

1. 計算資源、モデル、データのスケーリング
2. アルゴリズム的パラダイム転換
3. 再帰的自己改善
4. マルチエージェント協調と集団的主体性

これらの経路を見ると、AGI が実現した後も AI の能力向上が続き、やがて ASI へ至る可能性があるように思われます。

しかし、その進歩が順調に続くとは限りません。

AI の性能を高めるには、データ、計算資源、電力、半導体、研究成果、物理的実験、社会的受容など、多数の条件が必要です。

そのどれか一つが不足しても、進歩は遅くなる可能性があります。

複数の制約が同時に作用すれば、AI の能力向上が長期間停滞するかもしれません。

原報告書の Section 5.5 は、こうした制約を frictions and bottlenecks——摩擦とボトルネックとして整理しています。

ここでいう「摩擦」と「ボトルネック」は、同じではありません。

摩擦

進歩を遅らせるが、資源投入、技術革新、制度変更などによって克服できる制約です。

ボトルネック

能力向上を強く制限し、その問題を解決しない限り、次の段階へ進めなくなる制約です。

ハード・ブロッカー

現在の技術路線では原理的または実質的に克服できず、進歩を長期間停止させる障壁です。

原報告書は、現時点では、どの制約が単なる摩擦にとどまり、どれが進歩を止めるハード・ブロッカーになるかは分からないと強調しています。

そのため、著者たちは未来を断定するのではなく、各制約の影響を検証すべき未解決の研究課題として提示しています。（原報告書、Section 5.5）

第1節 六つの主要なボトルネック

原報告書の Table 4 は、AGI から ASI への移行を妨げ得る主要な要因を、次の六つに整理しています。

1. Data Wall——データの壁
2. Economic and Natural Resource Demand——経済的・自然資源需要の急増
3. Neural Paradigm Is Insufficient——現在のニューラル・パラダイムの限界
4. Research Gets Harder——研究の難化
5. Abstraction Barrier——抽象化の壁
6. Deliberate Slowdown——意図的な減速

抽象化の壁からは、さらに重要な問題として、

Embodied Bottleneck——身体性・現実接地のボトルネック

が導かれます。

これらは互いに独立ではありません。

たとえば、データの壁を越えるために物理世界で大量の実験を行えば、エネルギー、設備、時間の制約が強くなります。

AI が研究を自動化して研究の難化を相殺しても、実験設備の建設や薬剤の臨床試験はデジタル速度では進みません。

したがって、AGI から ASI への速度は、最も進歩した一要素ではなく、最も遅い不可欠な工程によって決まる可能性があります。

第2節 第一のボトルネック——データの壁

2.1 なぜデータが不足するのか

現在の基盤モデルは、大量の文章、画像、音声、映像、コードなどから学習します。

モデル規模と計算量を増やし続けるなら、それに対応する量と質の学習データも増やさなければなりません。

しかし、人間が新たに生み出す高品質な情報の量は、AIモデルの計算規模ほど急速には増えません。

原報告書は、モデル規模の成長が、意味ある新規テキストが世界で生産される速度を上回っていると指摘しています。

画像、音声、映像にはテキストより長い余地がある可能性がありますが、それらも人間だけによる生産では、計算需要と同じ速度では増えないかもしれません。（原報告書、Section 5.5）

2.2 データ量とデータ品質

データの壁は、単にファイル数が足りないという問題ではありません。

重要なのは、

- 正確である
- 重複が少ない
- 多様である
- 学習対象より少し難しい
- 現実世界と対応している
- 結果を検証できる

データです。

インターネット上には巨大な情報がありますが、同じ文章の複製、誤情報、自動生成された低品質コンテンツも含まれます。

データを増やしても、信頼性や新規性が低ければ、能力向上につながらない可能性があります。

2.3 AI 生成データは解決になるか

生成 AI は、文章、画像、音声、映像を大量に作れます。

そのため、AI に次世代 AI の訓練データを作らせることが考えられます。

しかし、現在のモデルが生成したデータを、検証せずに次のモデルへ反復学習させると、

- 誤りが蓄積する
- 出力の多様性が失われる
- 既存の偏りが増幅される
- 能力が頭打ちになる
- モデルが劣化する

危険があります。

原報告書は、単純な自己生成データの反復利用では、性能停滞や退化が起こる可能性を認めています。（原報告書、Section 5.5）

2.4 高品質な合成データ

ただし、AI 生成データがすべて無効というわけではありません。

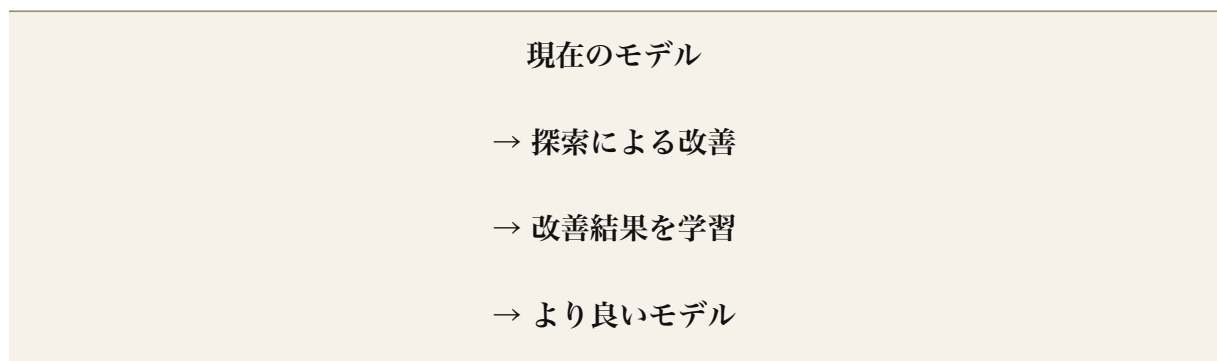
次のような仕組みがあれば、基盤モデル単独の能力を超える高品質データを作れる可能性があります。

- 推論時に多数の候補を探索する
- コード実行によって正しさを確認する

- 数学的検証器を使う
- シミュレーションで結果を評価する
- 自己対戦を行う
- 外部環境からフィードバックを得る
- 複数の AI による独立評価を行う

たとえば AlphaZero では、現在のモデルを使って探索を行い、モデル単独より良い手を見つけ、その結果を次の学習へ戻します。

これにより、



という循環が成立します。

原報告書は、推論時計算、検索、自己対戦、強化学習、シミュレーション、エージェントと環境の相互作用が、データの壁を越える主要候補になるとしています。（原報告書、Section 5.5）

2.5 数百万・数十億人の利用者が生むデータ

高度な AI が世界中で利用されれば、利用者と AI の相互作用そのものが大量の学習資源になります。

人間が AI の回答を、

- 採用する
- 修正する

- 拒否する
- 実行結果を返す
- さらに質問する

ことで、AI の能力境界付近の情報が生まれます。

特に、各応答について推論時計算を多く使い、その結果を選別・蒸留できれば、人間生成コンテンツとは異なる新しいデータ循環が生じます。

したがって、データの壁は、スケーリングを遅らせる摩擦にはなっても、必ずしも決定的障壁になるとは限りません。

2.6 データの壁が強く作用する経路

スケーリング経路

最も直接的に影響します。

より大きなモデルを訓練しても、十分な新規データがなければ性能向上が鈍化します。

パラダイム転換

データ効率の高い学習方式が生まれれば、壁を弱められます。

再帰的自己改善

AI が高品質データを生成・評価できれば、自己改善を加速できます。

しかし、誤った自己生成データが循環すれば、改善ではなく劣化を起こします。

マルチエージェント経路

多数の AI が相互作用することで新しいデータを生成できます。

ただし、同じ基盤モデル同士の会話では、既存の誤りや概念体系を反復するだけになる可能性もあります。

第3節 第二のボトルネック——経済的・自然資源需要の急増

3.1 AIは純粋な情報技術ではない

AIはデジタル技術ですが、その実行には大規模な物理基盤が必要です。

- 半導体
- データセンター
- 電力
- 冷却設備
- 通信網
- 水
- 土地
- 銅や希少鉱物
- 製造装置
- サプライチェーン
- 巨額の投資

です。

モデル規模、推論時計算、AI個体数を何桁も増やすなら、こうした資源も拡大しなければなりません。

原報告書は、現在のAIパラダイムを継続的にスケールするために必要な経済的、技術的、自然的資源の増加を維持できない可能性を、第二の主要ボトルネックとしています。（原報告書、Section 5.5）

3.2 経済的に持続できるか

AI開発費用が高くても、それ以上の経済的価値を生み出すなら投資は続きます。

AI が、

- ソフトウェア開発を自動化する
- 医薬品開発を高速化する
- 企業経営を改善する
- 知的労働を代替する
- 科学技術の進歩を速める

なら、その利益を次の AI 設備へ再投資できます。

したがって、AI スケーリングの限界は、費用の絶対額だけでは決まりません。

重要なのは、

追加投資による経済的収益が、追加費用を上回るか。

という関係です。

AI による経済成長が急速なら、巨大な計算基盤を維持できる可能性があります。

反対に、計算量を 10 倍にしても能力や収益がわずかしか増えなければ、投資は鈍化します。

(原報告書、Section 5.5)

3.3 半導体だけではない

計算資源を拡大するには、チップの製造量だけでなく、チップ間の通信も重要です。

原報告書は、物理的・工学的制約として、

- メモリ帯域
- チップ間通信
- データ移動
- 大規模並列処理の効率

を挙げています。

多数のチップを接続しても、必要な情報を高速に交換できなければ、計算資源を有効に利用できません。

モデルが巨大になるほど、計算そのものより、メモリや他のチップとの間でデータを移動する時間が主要な制約になる可能性があります。（原報告書、Section 5.5）

3.4 AIが資源効率を改善する可能性

この制約には、強力な反作用もあります。

AI自身が、

- 半導体設計を改善する
- アルゴリズム効率を高める
- データセンターを最適化する
- 冷却や電力利用を改善する
- 新しい材料を発見する
- エネルギー生産を効率化する

可能性です。

同じ能力を10分の1の電力や計算量で実現できれば、ハードウェアを増設したのと同じ効果があります。

したがって、重要なのは物理的計算資源だけでなく、実効計算量です。

3.5 資源需要が強く作用する経路

スケーリング経路

最も大きな制約になります。

能力向上が計算量の増加に強く依存するほど、経済・資源問題は重大になります。

パラダイム転換

新しいアルゴリズムやハードウェアが、同じ能力を少ない資源で実現すれば、制約を大幅に緩和できます。

再帰的自己改善

AIがアルゴリズム、チップ、エネルギー効率を改善すれば、資源の壁を押し広げられます。

ただし、工場や発電所の建設には現実時間が必要です。

マルチエージェント経路

AI 個体数の増加は計算需要を直接増やします。

集団知能の向上が劣線形なら、個体を増やす費用が成果を上回る可能性があります。

第4節 第三のボトルネック——現在のニューラル・パラダイムの限界

4.1 現在の方法で AGI に到達できるのか

今日の主流 AI は、おおむね次の要素から構成されています。

- 大規模ニューラルネットワーク
- Transformer
- 大量データによる事前学習
- 確率的勾配降下法
- 追加学習・事後訓練
- 推論時スケーリング
- 検索やツール利用
- エージェント型の足場

原報告書は、これらを組み合わせた現在のパラダイムで、AGI、さらに ASI に到達できるかは未解決だとしています。（原報告書、Section 5.5）

4.2 現在のパラダイムが進化する可能性

Transformer の事前学習だけを見れば、明らかな不足があります。

しかし現在のシステムはすでに、

- 高度な事後訓練
- 長い推論
- 検索
- 外部ツール
- エージェント機能

- 記憶
- 自己評価

を取り込みつつあります。

そのため、現在の枠組みが多数の追加要素を取り込みながら、比較的連続的に AGI へ近づく可能性があります。

この場合、現在のパラダイムの限界はハード・ブロッカーではなく、技術的摩擦にとどまります。

4.3 根本的問題かもしれない欠点

一方、原報告書は、現在観測されている次の問題の一部が、事前収集データを予測する学習方式そのものに関係している可能性を挙げています。

- 幻覚
- プロンプトインジェクションへの脆弱性
- 不確実性を適切に表現できない問題
- リスクや曖昧さを考慮した意思決定の弱さ
- 第三者データから行動を学ぶことで生じる自己欺瞞
- 新しい概念形成の限界

これらが規模拡大や追加訓練で継続的に改善するのか、ある段階で根本的限界として現れるのかは分かっていません。（原報告書、Section 5.5）

4.4 この壁への対抗要因

- 新しいアーキテクチャ
- 継続学習
- 世界モデル

- 強化学習中心の事前学習
- 現実環境との能動的相互作用
- 新しい最適化原理
- ニューロモーフィック計算
- アナログ計算

などが考えられます。

さらに、人間レベルに達していない AI でも、AI 研究を支援することで、新しいパラダイムの発見を加速できる可能性があります。

4.5 四経路との関係

スケーリング経路

現在のパラダイムに本質的な天井があれば、規模拡大だけでは ASI に到達できません。

パラダイム転換

このボトルネックを直接克服する経路です。

再帰的自己改善

AI が新しい学習原理やアーキテクチャを発見できれば、壁を突破する可能性があります。

マルチエージェント経路

個体が人間レベル付近で停滞しても、多数の AI を組織すれば、集合的には人間組織を超えられるかもしれません。

第5節 第四のボトルネック——研究の難化

5.1 「低い枝の果実」がなくなる

研究分野の初期には、比較的少ない努力で大きな改善を見つけられることがあります。

しかし、明白な方法が試されるにつれ、残る問題は難しくなります。

- より多くの研究者が必要
- より大規模な実験が必要
- より高価な設備が必要
- より長い探索が必要
- 仮説空間が巨大になる

可能性があります。

原報告書は、分野が成熟するほど研究者一人当たりの生産性が低下し、一定の技術進歩を維持するために、より多くの資源が必要になるという「研究が難しくなる」仮説を取り上げています。（原報告書、Section 5.5）

5.2 AI 研究にも同じことが起きるか

AI 開発の初期には、比較的単純な改善でも大きな成果が得られたかもしれません。

しかし将来は、

- 小さな性能向上に巨大な実験が必要
- 新しいアルゴリズムの発見確率が低下
- 安全性や信頼性の検証が複雑化
- 多数の能力間のトレードオフが増加

する可能性があります。

この場合、AI が研究を支援しても、課題の難化と研究能力の向上が相殺し合います。

5.3 人工研究者が状況を逆転させる可能性

一方、人間研究者を増やすには、教育に長い年月と費用が必要です。

人工研究者なら、計算資源があれば比較的短期間に多数稼働できる可能性があります。

数十万または数百万の AI 研究者を並列に動かせるなら、

- 仮説を大量生成する
- 多数の実験を設計する
- 文献全体を検索する
- 結果を自動分析する
- 異なる研究経路を並行して試す

ことができます。

原報告書は、研究が難化するとしても、人工研究者の大量投入による生産性向上の方が大きければ、この摩擦は限定的になる可能性がある」と論じています。（原報告書、Section 5.5)

5.4 認知労働と実験を区別する

ここで重要なのは、研究の認知的部分と物理的部分を分けることです。

AI は、

- 文献調査
- 数学的推論
- 仮説形成
- コード作成
- データ分析

を高速化できます。

しかし、

- 化学反応を待つ
- 生物を育てる
- 臨床試験を行う
- 宇宙を観測する
- 新素材を製造する

には現実時間が必要です。

認知的研究が完全に自動化されても、実験や観測が新しい制約になる可能性があります。

5.5 四経路との関係

スケーリング経路

能力向上のための実験規模が急増すれば、限界効用が低下します。

パラダイム転換

新しいアイデアを見つけること自体が難しくなれば、転換の頻度が下がります。

再帰的自己改善

人工研究者の大量投入は研究難化を相殺できます。

しかし、一世代ごとの改善に必要な資源がさらに速く増えれば、自己改善は頭打ちになります。

マルチエージェント経路

多数の研究 AI による分業は有効ですが、課題が並列化しにくい場合や、統合・検証が難しい場合には効果が限定されます。

第6節 第五のボトルネック——抽象化の壁

6.1 人間の概念を学ぶ AI

現在の AI は主として、人間が作った文章、画像、コード、分類、理論から学びます。

これらは、生の世界そのものではありません。

人間が長い歴史の中で、現実を観察し、そこから抽出した概念を、言語や記号へ変換したものです。

たとえば、

- ・ 力
- ・ 原因
- ・ 遺伝子
- ・ 市場
- ・ 重力
- ・ 民主主義
- ・ 無限
- ・ 確率

といった概念は、世界にそのまま文字で書かれているわけではありません。

人間が観測、実験、思考、文化的対話を通じて形成した抽象概念です。

原報告書の「抽象化の壁」とは、主として人間の認知的産物から学ぶ AI が、既存の人間概念の範囲内に閉じ込められる可能性を指します。（原報告書、Section 5.5）

6.2 再結合と新概念の発見は異なる

現在の AI は、既存の概念を大量に学び、それらを新しい形で組み合わせることに優れています。

しかし、

人間がまだ持っていない概念を、生の現実から新しく形成できるのか。

は別の問題です。

原報告書は、仮に現代のモデルへ膨大なデータを与えても、内容がニュートン以前の科学知識に制限されていたなら、その AI が独力で一般相対性理論や量子力学に到達できるかを問いかけます。

微積分、重力、電磁気などの概念的基盤が存在しない状態から、現在の予測学習だけで新しい概念体系を構成できるかは不明です。（原報告書、Section 5.5）

6.3 ベンチマーク飽和の別の解釈

AI が多くの試験で人間水準へ近づいていることは、知能全体が同じ速度で人間を超えることを意味するでしょうか。

抽象化の壁の仮説から見ると、現在の急速な性能向上は、

- ・ 人間が定義した問題
- ・ 人間が検証できる正解
- ・ 人間が作った概念
- ・ 人間の言語で表現された課題

の内部での習熟かもしれません。

つまり、AI は人間の知識空間の中では急速に向上していても、その知識空間そのものを超える速度は別問題です。

6.4 個体 AGI の天井と集合的 ASI

抽象化の壁が強ければ、一体の AI が人間の概念形成能力を大幅に超えることは難しいかもしれません。

しかし、原報告書は、それでも集合的 ASI が成立する可能性を指摘しています。

人間レベルの AI を大量に稼働させれば、

- 高速な検索
- 巨大な記憶
- 多数の仮説の並列検討
- 大規模な専門分業

によって、人間集団をはるかに超える成果を出せる可能性があるためです。

したがって、抽象化の壁が一体の AI の質的能力を制限しても、量的・組織的拡大による ASI を完全には排除しません。（原報告書、Section 5.5）

6.5 抽象化の壁を越える方法

原報告書が示唆するのは、AI を人間生成データだけで訓練するのではなく、

- 生のセンサーデータ
- 物理環境との相互作用
- ロボットによる行動
- 自律的実験
- 強化学習
- 能動的仮説検証

を通じて、世界から直接概念を形成させる方法です。

これは現在の事前学習中心の方式から、より対話的・経験的な学習へのパラダイム転換を必要とする可能性があります。（原報告書、Section 5.5）

第7節 身体性・現実接地のボトルネック

7.1 新概念は現実によって検証されなければならない

AIが新しい物理法則や生物学的機構をデジタル速度で提案できても、それが正しいかどうかは、現実世界で確認しなければなりません。

原報告書は、新しい概念と、その概念を使った予測規則を現実にも照らして検証する必要性を、Embodied Bottleneck——身体性・現実接地のボトルネックと呼んでいます。（原報告書、Section 5.5）

ここでいう「身体性」は、人型ロボットの身体だけを意味しません。

AIが、

- ・ センサーを通じて観測する
- ・ 装置を操作する
- ・ 実験する
- ・ 結果を測定する
- ・ 仮説を現実にも照らして修正する

という、物理世界との循環を持つことです。

7.2 物理時間は圧縮できない

一部の現象には、固有の時間尺度があります。

- ・ 化学反応
- ・ 生物の成長
- ・ 疾患の進行
- ・ 気候変動
- ・ 材料の劣化

- 人間社会の長期反応

です。

AIの思考速度を千倍にしても、植物が千倍速く成長するわけではありません。

臨床試験の長期的副作用を数秒で直接観測することもできません。

シミュレーションで一部を代替できますが、複雑な生物、気候、社会を、長期間にわたり完全な精度で再現することは難しい場合があります。（原報告書、Section 5.5）

7.3 知能爆発への直線的な制約

再帰的自己改善がデジタル領域だけで完結するなら、改善サイクルは非常に短くなる可能性があります。

しかし、新しい知識の獲得に物理的実験が不可欠なら、

AIの能力向上速度は計算機の手速ではなく、経験科学の手速に制約される

可能性があります。

これは知能爆発の見通しを大きく変えます。

AIが数分で一万個の仮説を作れても、信頼できる実験結果を得るのに数か月必要なら、研究サイクル全体は数か月より短くなりません。

7.4 デジタル領域と物理領域の差

このボトルネックの影響は分野によって異なります。

デジタル領域

- 数学
- ソフトウェア
- アルゴリズム

- 一部の形式科学

では、自動評価や高速な検証が可能です。

自己改善は比較的速く進む可能性があります。

物理・生物・社会領域

- 材料科学
- 医学
- ロボット工学
- 気候科学
- 社会制度

では、現実世界の観測や介入が必要です。

進歩は物理的・制度的時間に制約されます。

したがって、最初の ASI はすべての分野で均一に超人的になるのではなく、デジタルに検証可能な分野から先行する可能性があります。

第8節 第六のボトルネック——意図的な減速

8.1 技術的に可能でも、社会が許可しない場合

これまでのボトルネックは、主として技術、資源、知識に関するものでした。

しかし、AIの進歩は社会の中で起こります。

技術的・経済的にASIへ進める場合でも、人間社会が意図的に速度を下げる可能性があります。

原報告書は、その要因として、

- ・ 悪意ある利用
- ・ 重大事故
- ・ 制御喪失の危険
- ・ 軍事的・政治的濫用
- ・ 社会文化的損害
- ・ 雇用や格差への影響
- ・ 社会的反発

を挙げています。（原報告書、Section 5.5）

8.2 労働と社会契約

AGIが多くの認知労働を、人間より安価に提供できるようになれば、経済の中心的資源が労働から資本へ移る可能性があります。

AIと計算資源を所有する者へ所得と権力が集中すれば、

- ・ 雇用
- ・ 所得分配
- ・ 税制

- 社会保障
- 人間の社会参加
- 尊厳と自律性

を支えてきた制度を再設計する必要があります。

社会がこの変化に対応できなければ、AI 開発への政治的反発が強まり、研究や配備の制限につながる可能性があります。（原報告書、Section 5.5）

8.3 重大事故と社会的許可

高度な AI が、金融、医療、電力、軍事、行政などの重要インフラへ組み込まれると、一つの事故が広範囲へ影響する可能性があります。

大規模な事故や信頼できるニアミスが起きれば、

- 世論
- 法的責任
- 保険制度
- 許認可
- 投資判断

が変化し、次の能力拡大が政治的・商業的に困難になるかもしれません。

原報告書は、複雑で密結合した AI システムでは、個別の技術故障だけでなく、組織的意思決定から重大事故が生じ得る点も指摘しています。（原報告書、Section 5.5）

8.4 ガバナンスは単なるブレーキではない

規制やガバナンスは、進歩を止めるだけのものではありません。

- 能力基準による許認可
- 配備前評価

- 事故報告
- セキュリティ基準
- 段階的な能力解放
- 独立監査

を通じて、技術開発の方向と質を変えることもできます。

適切なガバナンスによって重大事故を防げれば、長期的には社会的信頼を維持し、持続可能な技術発展を可能にするかもしれません。（原報告書、Section 5.5）

8.5 国際競争による反作用

意図的な減速には強い反作用があります。

一つの企業や国家が開発を遅らせても、他の主体が経済的・軍事的優位を求めて進める可能性があります。

- 国際協調の困難
- 規制の弱い地域への開発移転
- 安全性より競争速度を重視する圧力
- 軍事・経済的優位への懸念

により、単独の減速政策は持続しにくいかもしれません。

原報告書は、国家間競争や市場圧力が減速要求を上回る可能性を認めています。（原報告書、Section 5.5）

第9節 六つの制約はどのように重なり合うのか

9.1 データの壁と抽象化の壁

合成データによってデータ量を増やせても、そのデータが人間の既存概念を再構成したものにすぎなければ、抽象化の壁は越えられません。

量の不足を解決しても、質的に新しい概念形成の不足が残ります。

9.2 抽象化の壁と身体性

新しい概念を形成するには、生の世界との相互作用が必要になるかもしれません。

すると、センサー、ロボット、実験設備、観測時間が新たな制約になります。

デジタル知能の高速性が、そのまま科学的発見速度へ変換されなくなります。

9.3 研究難化と再帰的自己改善

AI 研究者を大量投入すれば、研究難化を相殺できます。

しかし、次の改善に必要な実験や計算量が、人工研究者の増加より速く増えれば、自己改善は鈍化します。

重要なのは、

研究能力の増加率と、研究課題の難化率のどちらが大きいのか

です。

9.4 資源制約と国際競争

資源が不足するほど、半導体、電力、鉱物、データセンターをめぐる競争が強まります。

その競争は能力向上を促進する一方、軍事化、集中、社会的反発を生み、意図的な規制を招く可能性があります。

9.5 事故と進歩速度

急速な再帰的改善やマルチエージェント化は、社会へ十分に検証されていないシステムを導入する圧力を高めます。

その結果、重大事故が起これば、今度は強い規制と減速が生じます。

したがって、技術加速と社会的減速は、互いに独立ではなく、フィードバック関係にあります。

第10節 AGIでちょうど進歩が止まる可能性

原報告書は、人間レベル AGI が実現したと仮定した場合、AI の進歩が正確に人間レベルで停止する可能性は低いと考えています。

一体の AI が抽象化の壁などによって人間レベル付近で停滞しても、

- より多くの計算資源を使う
- より高速に稼働させる
- 多数の AGI を複製する
- 集団として組織する

ことで、集合的能力を高められる可能性があるためです。

大規模な人間レベル AGI 集団が、人間の専門家組織を超える可能性は高いと原報告書は推測しています。

したがって、AI の進歩が AGI で完全停止するには、複数の摩擦が同時に強いハード・ブロッカーとして作用しなければならない可能性があります。（原報告書、Section 6）

10.1 AGI 以前に停滞する可能性

現在のパラダイムが、継続学習、世界理解、長期計画、信頼性などの問題を解決できないなら、AGI 到達前に進歩が停滞する可能性があります。

この場合、ASI 以前に、AGI そのものが新しいパラダイムを待つことになります。

10.2 AGI から弱い ASI へ比較的滑らかに進む可能性

一度、人間レベルで汎用的に働ける AI が実現すれば、

- 個体数の増加
- 動作速度の向上
- 並列化
- 分業
- AI 研究支援

によって、弱い ASI へ比較的連続的に進む可能性があります。

原報告書は大きな不確実性を認めつつ、AI が正確に人間レベルで停滞するよりも、AGI 以前に停滞するか、AGI から弱い ASI へ進む方が相対的にあり得ると述べています。（原報告書、Section 6）

10.3 急激な知能爆発

再帰的自己改善が摩擦を上回れば、AGI から ASI への移行は短期間に起こる可能性があります。

しかし身体性、物理実験、製造、資源などが強く作用すれば、知能爆発は限定的になるかもしれません。

したがって、重要なのは「知能爆発は起きるか」という一問だけではありません。

- どの分野で起きるか
- どの能力が先行するか
- どの工程で物理世界に制約されるか
- デジタル領域と現実領域の速度差はどの程度か

を区別する必要があります。

ERF Commentary

第11節 大学はAIの摩擦を減らすだけの制度ではない

大学は、AI開発を加速する研究機関として期待されます。

- 新しいアルゴリズム
- 高効率なハードウェア
- 高品質データ
- AI 安全性
- ロボット技術
- 科学的発見

を通じて、いくつかのボトルネックを解消できます。

しかし、大学の役割を、AI 進歩の摩擦を取り除くことだけに限定してはなりません。

一部の摩擦は、人間社会を守るために必要だからです。

たとえば、

- 倫理審査
- 再現性確認
- 安全性評価
- 社会的対話
- 人権保護
- 環境影響評価

は、技術開発を遅くすることがあります。

しかし、その遅さは非効率ではなく、重大な誤りを防ぐための制度的慎重さです。

第12節 データの壁と大学の知識資源

大学は、論文、実験データ、図書館、研究記録など、大量の高品質な知識資源を保有しています。

AI時代には、それらが貴重な学習資源になります。

しかし、大学の知識を無制限に AI 企業へ提供するだけでは問題があります。

- 著作権
- 研究参加者のプライバシー
- 機密性
- 学術的帰属
- 研究者の利益
- 公共財としての知識アクセス

を考える必要があります。

大学は、学術データを閉鎖するだけでも、無条件に開放するだけでもなく、公共的・倫理的なデータガバナンスを形成する役割を持ちます。

第13節 抽象化の壁と人間の研究

AI が人間の既存知識を高速に再構成できても、新しい概念を作る能力が限定されるなら、人間研究者の役割は依然として重要です。

特に、

- 現実の違和感に気付く
- 既存分類そのものを疑う
- 異分野の概念を移植する
- 新しい問いを作る
- 言葉になる前の経験を概念化する

能力です。

大学教育は、AI が知識を提供できる時代ほど、既存知識を覚えるだけでなく、概念を形成し直す力を育てなければなりません。

第14節 身体性と大学

知識は、すべてテキストやデータベースの中にあるわけではありません。

- ・ 実験室で装置を扱う
- ・ 患者と向き合う
- ・ フィールド調査を行う
- ・ 教室で他者と議論する
- ・ 地域社会へ参加する
- ・ 芸術を制作する

経験から形成される知識があります。

AI時代の大学がオンライン情報処理だけに傾けば、こうした身体的・社会的・状況的な学習を失う危険があります。

原報告書の身体性ボトルネックは、AIにとって現実接地が必要であるだけでなく、人間教育にとっても、世界との直接的関係が不可欠であることを示唆します。

第15節 研究の速度と学問の時間

AIによって研究サイクルが高速化しても、学問には時間を必要とする側面があります。

- ・ 結果を再現する
- ・ 長期的影響を観察する
- ・ 異なる立場から批判する
- ・ 理論の意味を理解する
- ・ 社会が知識を受け入れる
- ・ 世代を超えて評価する

ためです。

速く生成された知識が、速く理解されるとは限りません。

大学は、AI時代に研究速度を上げるだけでなく、知識を熟成させ、批判し、位置付けるための時間を守る必要があります。

第16節 意図的な減速と大学の公共的責任

大学は、技術の進歩を無条件に支持する組織ではありません。

また、技術を恐れて停止させるだけの組織でもありません。

大学には、

- どの進歩が望ましいか
- どのリスクを許容できるか
- 誰が利益を得るか
- 誰が負担を負うか
- どの条件で速度を下げるべきか

を公共的に議論する役割があります。

適切な減速は、反技術的な行為ではありません。

社会が技術を理解し、制度を整え、人間の自律性と尊厳を守るための時間を確保する行為です。

第17節 大学は「アクセル」と「ブレーキ」の両方を担う

AI時代の大学は、一方では、

- 新しい知識
- 新しい技術
- 新しい人材
- 新しい制度設計

を生み出し、社会変革を促進します。

他方では、

- 批判
- 検証
- 倫理
- 公共的熟議
- 長期的視点

を通じて、無制限な加速を抑制します。

この二つは矛盾ではありません。

安全で公共的に正当な進歩を実現するためには、加速と慎重さの両方が必要です。

第18節 Ecological University との関係

Ronald Barnett の Ecological University は、大学を複数の生態系へ責任を負う制度として捉えます。

AI の発展も、

- 知識の生態系
- 経済の生態系
- 自然環境
- 政治制度
- 労働と生活
- 人間の人格形成
- 国際秩序

へ同時に作用します。

AI 能力だけを最大化すれば、他の生態系を損なう可能性があります。

計算資源を拡大すれば、環境負荷が増えるかもしれません。

労働を自動化すれば、生産性は高まっても、社会参加や所得配分が損なわれるかもしれません。

科学を高速化すれば、研究成果は増えても、人間が理解し統治する能力を超えるかもしれません。

Ecological University の役割は、単一の技術指標を最大化するのではなく、複数の生態系間の関係を見渡し、責任ある均衡を探ることにあります。

第19節 ボトルネックを取り除くことが常に善とは限らない

データ、電力、半導体、研究効率の壁は、一般には解決すべき技術的制約と見なされます。

しかし、すべての制約を最速で取り除くことが、社会的に望ましいとは限りません。

たとえば、技術的に危険なシステムの能力拡大を妨げる制約は、安全性が確立するまで有益かもしれません。

問題は、制約が存在するか否かではなく、

どの制約を克服し、どの制約を安全装置として維持すべきか。

です。

これは工学だけでなく、民主的判断と公共哲学の問題です。

結論

『From AGI to ASI』の Section 5.5 は、AGI から ASI への移行を妨げ得る六つの主要な摩擦とボトルネックを示しています。

1. 高品質な学習資源が不足するデータの壁
2. 投資、半導体、電力、鉱物、設備を拡大できない経済的・自然資源の壁
3. 大規模事前学習と勾配降下法を中心とする現在のニューラル・パラダイムの限界
4. 発見が進むほど次の発見に多くの資源を要する研究の難化
5. 人間の概念体系を超える新しい概念を形成できない可能性としての抽象化の壁
6. 事故、リスク、社会的反発、規制による意図的な減速

さらに抽象化の壁から、AI が新しい概念を現実世界で検証しなければならない身体性・現実接地のボトルネックが生じます。（原報告書、Section 5.5）

これらの制約のどれも、現時点で決定的なハード・ブロッカーだとは証明されていません。

データの壁には、合成データ、探索、自己対戦、シミュレーションがあります。

資源の壁には、アルゴリズム効率、半導体革新、AI による技術開発があります。

研究の難化には、人工研究者の大量投入があります。

個体 AI の限界には、マルチエージェント集合知があります。

しかし、反対に、これらの対抗手段が十分に機能する保証もありません。

したがって、原報告書の最も重要な主張は、

AGI から ASI へ進めるかどうかは、一つの技術的突破によってではなく、加速要因と減速要因の相互作用によって決まる

という点です。

AI の未来は、単純な直線でも、必然的な知能爆発でもありません。

データ、計算、研究、物理世界、社会制度、政治的判断が相互に影響する、複雑な社会技術的過程です。

そして高等教育に求められるのは、AI 発展を単に促進または阻止することではありません。

どの摩擦が不必要な障害で、どの摩擦が社会を守る安全装置なのかを見極め、技術的可能性を人間的・公共的・生態学的目的へ結び付けることです。

第9回では、原報告書の Section 6 と Section 7 へ進み、四つの経路と六つのボトルネックを統合した全体像、今後必要となる研究課題、そして Google DeepMind が最終的にどのような未来認識を示しているのかを整理します。