

AGI への岐路：スケーリング仮説と新アーキテクチャ探求の深層分析

Gemini

エグゼクティブサマリー

本レポートは、現代の人工知能（AI）における中心的な議論、すなわち「現在の AI 技術の延長線上にある継続的な発展だけで汎用人工知能（AGI）は達成可能なのか、それとも根本的なパラダイムシフトが必要なのか」という問いについて、包括的かつ深層的な分析を提供する。この議論は、AI の未来、その社会的影響、そして数十億ドル規模の投資の方向性を決定づけるものであり、技術戦略家、研究者、投資家にとって極めて重要な意味を持つ。

中心的論点： AGI への道は、現在のトレンドの単純な延長でも、それからの完全な離脱でもない。むしろ、それは**収斂進化（convergent evolution）**のプロセスである可能性が高い。このプロセスにおいて、現行モデルの圧倒的なスケールは、スケーリングのみのパラダイムに内在する限界を克服するために、世界モデル、因果的推論、そして身体性（エンボディメント）に関連する根源的なアーキテクチャ革新と統合される必要に迫られるだろう。

主要な分析結果：

- スケーリング仮説の限界：** 経験的に強力な成果を上げてきた「スケーリング仮説」は、現在、収穫逨減の兆候を見せており、真の推論能力や汎化能力に関して根本的な理論的限界に直面している。計算資源、高品質データ、エネルギー消費といった物理的制約も、このアプローチの持続可能性に疑問を投げかけている。
- 「AGI」定義の戦略的変遷：** 「AGI」という概念自体が、かつての「意識を持つ強い AI」という哲学的目標から、「経済的価値を持つタスクを遂行する能力」という実用的かつ経済的なベンチマークへと大きく移行している。この定義の再設定は、議論の枠組みそのものに大きな影響を与えており、現在のスケーリング中心のアプローチを正当化する側面を持つ。
- 代替アーキテクチャの必然性：** ヤン・ルカン氏が提唱する「世界モデル」や、ゲイリー・マーカス氏が主導する「ニューロシンボリック AI」といった代替アーキテクチャは、単なる学術的探求ではない。これらは、スケールアップされた大規模言語モデル（LLM）が露呈した具体的な失敗、すなわち因果関係の理解の欠如や信頼性の低さに対する直接的な応

答として登場した必然的なアプローチである。

4. 哲学的議論の現代的意義：「中国語の部屋」のような思考実験に代表される「理解」をめぐる哲学的議論は、現代においても極めて高い妥当性を持つ。それは、現在の AI が得意とする高度なパターンマッチングと、AGI に求められる真の認知との間に存在する質的なギャップを鋭く指摘しており、このギャップはスケールのみでは埋められない可能性を示唆している。

戦略的展望： AI 開発の次なるフェーズは、アプローチの「ハイブリッド化」によって特徴づけられるだろう。最も大きな進歩は、スケールされた LLM が持つ広範な知識ベースと、より構造化され、因果関係を扱え、物理世界に根差した推論メカニズムとを組み合わせたシステムから生まれる可能性が最も高い。したがって、将来の技術的優位性は、単に最大のモデルを構築する能力だけでなく、これらの異なるパラダイムをいかに巧みに統合できるかにかかっている。

第 1 章 目的地の定義：汎用人工知能（AGI）の変遷するゴールポスト

「現在の AI の発展だけで AGI は達成可能か」という問いに答えるためには、まずその目的地である「AGI」が何を指すのかを明確に定義する必要がある。しかし、この定義自体が固定されたものではなく、AI 研究の歴史の中で大きく変遷してきた。この定義の揺らぎを理解することは、現在の議論の全体像を把握するための不可欠な前提となる。

「強い AI」から経済的実用性へ

AI 研究の初期段階では、その目標はしばしば「強い AI（Strong AI）」という言葉で表現された。これは、哲学者ジョン・サールらが提起した概念であり、単に知的な振る舞いをするだけでなく、人間のような意識、自己認識、そして「真の理解」を持つ機械を指す、野心的な哲学的・認知科学的目標であった¹。この文脈における AGI は、計算能力だけでなく、内的な心的状態を持つ存在として構想されていた。

しかし、「意識」や「理解」といった内的状態は、外部から客観的に測定することが極めて困難である。この検証可能性の問題が、AGI の定義をより実用的で測定可能な方向へとシフトさせる大きな要因となった¹。研究の焦点は、検証不可能な内面を探索することから、観測可能な

外面の能力、すなわち「人間が実行できる広範な知的タスクを遂行する能力」へと移っていった⁴。この段階で重視されたのは、未知の課題に対応する柔軟性、自己改善能力、そして異なるドメイン間で知識を転移させる能力など、外部から評価可能な指標であった⁶。

この流れは、近年の大規模言語モデル（LLM）の台頭とともにさらに加速し、現代における AGI の定義を決定づけた。特に OpenAI が 2018 年に提示した憲章は、AGI を「ほとんどの経済的に価値のある仕事において人間を上回る、高度に自律的なシステム」と定義した¹。これは、評価軸を哲学的な問いから測定可能な汎用性へ、そして最終的には「経済的価値」へと移行させる画期的な再定義であった。この定義の下では、AGI はもはや抽象的な哲学概念ではなく、経済的リターンという明確な尺度で達成度が測られる、具体的な技術的目標となった。

この定義の変遷は、単なる学術的な整理ではない。それは、主要な AI 研究所による戦略的な行為と解釈できる。検証が困難な「強い AI」を目標とする限り、現在の LLM の能力では達成は程遠い。しかし、ゴールを「経済的に価値のある仕事の自動化」と再設定することで、現在のスケーリングという技術的アプローチの延長線上に、達成可能な目標として AGI を位置づけることができる。この再定義は、巨額の投資を正当化し、商業的インセンティブと研究開発の方向性を一致させる役割を果たしている。つまり、現在の技術で到達可能な場所にゴールポストを動かすことで、壮大なビジョンと現実的な開発ロードマップとの間の整合性を確保しているのである。

AGI の運用化：レベル分類とベンチマーク

定義が実用的なものへと移行する中で、AGI への進捗を段階的かつ客観的に評価するためのフレームワークが複数提案されている。その代表例が、研究者らによって提案された「AGI のレベル分類」である¹⁰。このフレームワークは、AI の能力を「性能（深さ）」と「汎用性（広さ）」という 2 つの次元で評価し、「レベル 1：新興（Emerging）」から「レベル 5：超人的（Superhuman）」といった段階に分類する。このような分類法は、「AGI か、否か」という二元論的な問いを、「現在どのレベルの AGI に到達しているか」という、より建設的でニュアンスのある議論へと転換させる。

さらに、AGI の能力を具体的に測るための概念的なベンチマークも考案されている。例えば、「コーヒー・テスト」は、平均的な家庭に入り、コーヒーメーカーを見つけ、豆を探し、コーヒーを淹れるといった一連のタスクを、人間の動きを模倣することなく自律的に実行できるかを問うものである⁵。同様に、「イケア・テスト」は、説明書を読んでイケアの家具を組み立てる能力を試すもので、言語理解、空間認識、物理的操作といった多様な能力の統合を要求する⁵。これらのテストは、単一のタスクにおける性能だけでなく、実世界における汎用的な問題解決能力を評価しようとする試みである。

残された曖昧さの問題

しかし、こうした運用化の努力にもかかわらず、「AGI」という用語は依然として科学的に厳密な定義を欠いている¹²。そもそも、基準となる「人間レベルの知能」自体が多面的で、単一の指標で測定することが難しい。知能の本質が何かという問いも、未だ解決には至っていない¹³。

この定義の曖昧さこそが、AGI 達成の可能性をめぐる議論を複雑にし、しばしば平行線をたどる原因となっている。ある者は経済的価値の創出をもって「AGI の火花」と評価し、またある者は真の理解や意識の欠如を指摘する¹。両者は異なる定義、異なる「地図」の上で議論しており、その評価が食い違うのは当然と言える。したがって、本レポートで扱う「AGI 達成は可能か」という問いは、どの定義を採用するかによってその答えが大きく変わりうるという、根本的な問題を内包していることを認識することが不可欠である。

第2章 「もっと多く」の道：スケーリング則、創発的能力、そして連続性の論拠

AGI は現在の AI 技術の延長線上で達成可能であるとする「連続性」の立場、すなわちスケーリング仮説は、近年の AI の驚異的な進歩を支えてきた最も強力なパラダイムである。この仮説は、モデルの規模を拡大し続けることで、いずれ人間レベルの知能が「創発」という経験則に基づいている。本章では、このスケーリング仮説を支える技術的基盤と、その支持者たちが描く AGI への道筋を詳細に分析する。

基礎：ニューラル・スケーリング則

スケーリング仮説の経験的根拠となっているのが、「ニューラル・スケーリング則 (Neural Scaling Laws)」である¹⁴。これは、AI モデルの性能（典型的には損失関数 (loss) の値で測定される）が、3 つの主要な要素、すなわちモデルサイズ（パラメータ数、）、データセットサイズ（トークン数、）、そして学習に使用する計算量（コンピュート、）を増大させるにつれて、べき乗則に従って予測可能に向上するという経験則である¹⁶。

2020 年に Kaplan らが発表した独創的な研究は、Transformer ベースの言語モデルにおいて、これらの要素と性能との間に、数桁のスケールにわたって非常に滑らかで予測可能な関係が存在することを示した¹⁶。この発見は、AI の性能向上がある種の「物理法則」に従うことを示唆し、AI 開発に明確な指針を与えた。すなわち、より優れた AI を作るためのレシピは、原理的には単純である。「より多くのデータで、より大きなモデルを、より多くの計算量を使って学習させる」ことだ。このスケーリング則の発見は、AI 開発を職人芸的な試行錯誤から、より予測可能で工学的な営みへと変貌させた。

現象：創発的能力

スケーリング則がもたらす最も興味深く、かつ重要な現象が「創発的能力 (Emergent Abilities)」である¹⁰。これは、小規模なモデルでは全く見られなかった能力が、モデルの規模がある特定の閾値を超えたときに、予期せず、しばしば急激に出現する現象を指す¹⁹。

例えば、多段階の論理的推論、文脈に応じた学習 (In-context Learning)、あるいは思考の連鎖 (Chain-of-Thought) プロンプティングといった高度な能力は、小規模なモデルでは観察されないが、GPT-3 のような巨大モデルで初めて顕著になった²⁰。これは、規模の「量的」な増加が、能力の「質的」な飛躍をもたらすことを示している。

スケーリング仮説の支持者たちは、この創発現象こそが AGI への道筋であると考えている。彼らの主張は、知能は明示的に設計されるものではなく、巨大なデータと計算の中で自己組織化し、創発するものであるというものだ。したがって、AGI を達成するための戦略は、スケーリングを続け、より強力で、より汎用的な能力が次々と創発してくるのを待つこと、ということになる¹⁰。

支持者たちのビジョン：Altman と OpenAI

このスケーリング哲学を最も強力に推進しているのが、OpenAI と CEO のサム・アルトマン氏である。彼のビジョンは、スケーリング仮説の論理的帰結を体現している。アルトマン氏は、AI モデルの知能は、学習と実行に用いられるリソース（計算量、データ）の対数にほぼ等しいと述べ、この進歩には「壁はない」と主張している²³。

このビジョンの下では、GPT-4 や GPT-5 といった各モデルのリリースは、AGI という最終目標に向けた連続的な道のりの一步一步に過ぎない。それぞれのバージョンは、前のバージョン

よりも大幅に優れていることが期待され、AGI は特定の完成形として定義されるのではなく、継続的な改善のプロセスそのものとして捉えられている²⁵。

究極の論拠：予測の帰結としての理解

スケーリング仮説を哲学的に支える最も強力な論拠の一つは、AI 研究の第一人者であるジェフリー・ヒントン氏によって提示された。彼の主張は、人間が生成した膨大なテキストコーパスにおいて、次に来る単語を高い精度で予測するというタスクを極めるためには、モデルはそのテキストが記述している世界の構造や意味、因果関係を「真に理解」せざるを得ない、というものである²⁷。

ヒントン氏は、人間が持つような主観的経験（クオリア）が理解の前提条件であるという考えを「くだらない（bullshit）」と一蹴する²⁷。彼によれば、モデルが示す驚異的な言語能力や推論能力そのものが、理解が存在する証拠である。この見解は、「AI は単なる統計的なパターンを模倣しているに過ぎない」という批判（いわゆる「確率的オウム（stochastic parrot）」論）に対する強力な反論となっている。予測という単純な目的関数を最大化するプロセスが、副産物として、あるいは必然的な結果として、世界についての深い理解を生み出すというのである。

しかし、この創発という現象は、スケーリング仮説の希望であると同時に、その最大の謎でもある。これらの能力がなぜ、どのようにして出現するのか、そのメカニズムは完全には解明されていない¹⁹。それが真の汎化能力の現れなのか、それとも極めて高度で複雑な記憶とパターンマッチングの産物なのかは、依然として議論の的である¹⁵。この「ブラックボックス」性は、AI の安全性という観点から重大な懸念を生じさせる。もし、我々が有益な能力の創発を予測できないのであれば、同様に、欺瞞、操作、あるいは意図しない目標の追求といった、有害な能力の創発も予測できないことになる¹⁹。AGI への希望をもたらすまさにその現象が、同時に最大のリスクの源泉ともなっているのである。

第3章 壁への到達：スケーリング・パラダイムの技術的・理論的境界

スケーリング仮説が AI の進歩を牽引してきた一方で、そのアプローチには無視できない限界が顕在化しつつある。これらの限界は、資源的な制約といった実用的な問題から、現行アーキテ

クチャに内在する根本的な理論的問題にまで及ぶ。本章では、スケーリングのみで AGI を目指すことへの包括的な反論を、具体的な証拠に基づいて展開する。

実用上の制約

スケーリング戦略は、その成功自体が自らの首を絞めるというパラドックスに直面している。指数関数的に増大する要求に対して、資源は有限である。

データの枯渇

LLM の性能は高品質な学習データに大きく依存するが、その供給源は無限ではない。インターネット上に存在する、利用可能な高品質のテキストや画像データは、すでに大規模なモデルの学習に大量に消費されており、その枯渇が現実的な問題として懸念されている¹⁷。合成データ（AI が生成したデータ）でこの問題を補う試みもあるが、品質や多様性の点で、人間が生成したデータに匹敵するかは未知数である。

計算資源とエネルギーコスト

モデルの規模が大きくなるにつれて、その学習と運用に必要な計算資源とエネルギーは天文学的に増大する。最新の GPU への需要は供給を上回り、その価格は高騰している²⁹。さらに、巨大なデータセンターが消費する電力は、環境への負荷という観点からも持続可能性が問われている¹⁷。地政学的な緊張が半導体の供給網に影響を与えるリスクも、この問題をさらに深刻化させている¹⁷。

収穫逡減

最も重要な点は、スケーリングによる性能向上のペースが鈍化し始めていること、すなわち「収穫逡減」の兆候が見られることである¹⁷。かつては、モデルサイズを倍にすれば性能が劇

的に向上したが、現在では同様の改善を得るためにより大きなスケールアップが必要になっている。DeepMind の「Chinchilla」モデルが示したように、単にモデルを巨大化するよりも、モデルサイズとデータ量のバランスを最適化の方が計算効率が良いことが明らかになった¹⁷。この「Chinchilla の変曲点」は、純粋なスケール至上主義の限界を露呈させた。OpenAI の元チーフサイエンティストであるイリヤ・サツキヴァー氏のような第一人者でさえ、事前学習をスケールアップさせることの利点が頭打ちになりつつあると指摘している²³。

理論的・アーキテクチャ上の限界

実用的な制約以上に深刻なのが、現在の LLM アーキテクチャ（主に Transformer ベース）が持つ根本的な限界である。

真の推論能力の欠如

LLM に対する最も根源的な批判は、それらが因果関係を理解して論理的に推論しているのではなく、あくまでデータ内の統計的な相関関係を学習しているに過ぎないという点である³⁰。ヤン・ルカン氏やゲイリー・マーカス氏といった批判者たちは、LLM を「確率的オウム」や「ヒューリスティクスの袋 (bags of heuristics)」と呼び、一貫性のある世界モデルを内蔵していないと指摘する³¹。この欠陥は、常識的な判断を要する場面で「とんでもない間違い (boneheaded mistakes)」を犯したり、一見もっともらしいが論理的に破綻した結論を導き出したりする原因となる³¹。

ハルシネーション（幻覚）問題

ハルシネーション、すなわち事実に基づかない情報を生成する問題は、単なる修正可能なバグではなく、LLM の確率的な性質に根差した本質的な欠陥である¹⁷。モデルは「事実」に基づいて応答を生成しているのではなく、「次に来る可能性が最も高い単語」を予測しているに過ぎない。そのため、特に学習データが少ない領域や、矛盾した情報が含まれる場合に、もっともらしい嘘を生成する傾向は避けられない。この信頼性の欠如は、医療や金融といった、正確性が絶対的に要求される高リスクな分野での応用を著しく制限する³¹。

専門家からの懐疑論

こうした限界は、AI 研究者の間でも広く認識されている。例えば、2025 年に行われた米国人工知能学会（AAAI）の調査では、回答した研究者の 76%が「現在の AI アプローチを単純にスケールアップするだけでは AGI の実現は見込めない」と考えていることが明らかになった³⁴。メタ社のヤン・ルカン氏や、Google Brain の共同設立者であるアンドリュー・ン氏など、業界を代表する重鎮たちも、現在の技術が抱える基本的な課題を指摘し、スケーリングのみに依存するアプローチに警鐘を鳴らしている²⁹。

これらの実用的および理論的な限界は、AI 業界に重要な転換点を強いている。それは、「より大きく（Bigger）」から「より賢く（Smarter）」への戦略的シフトである。単にモデルを巨大化させるという初期の戦略は、物理的・経済的な壁にぶつかり、同時にアーキテクチャ上の根本的な欠陥を解決できないことが明らかになった。その結果、AI 企業は、強化学習

（RLHF）や推論時計算（OpenAI の o1 モデルなど）といった「学習後改善」技術に多額の投資を行い、既存のモデルからより多くの知能を引き出そうと躍起になっている¹⁷。これは、生のスケールだけでは不十分であるという暗黙の承認に他ならない。この新しいフロンティア、すなわち「知能の効率性」の追求は、次章で詳述する代替アーキテクチャの探求へと直接つながっていく。一つのパラダイムの限界が、次なるパラダイムの触媒となっているのである。

第 4 章 アーキテクチャの分岐：AGI への代替経路を探る

スケーリング・パラダイムの限界が明らかになるにつれ、AI 研究の最前線では、AGI 達成に向けた根本的に異なるアプローチの探求が活発化している。これらのアプローチは、単なる既存モデルの改良ではなく、知能の本質に関する異なる仮説に基づいた、新しいアーキテクチャの構築を目指すものである。本章では、その中でも特に影響力のある 3 つの潮流——世界モデルと身体性、ニューロシンボリック統合、そして意識の科学的探求——を深掘りする。

4.1 世界モデルの指令と身体性

中核概念：ヤン・ルカンのビジョン

メタ社のチーフ AI サイエнтиストであるヤン・ルカン氏は、現在の LLM が AGI への道ではないと主張する最も著名な論者の一人である。彼の中心的な主張は、真の知能には「世界モデル（World Model）」が不可欠であるというものだ³²。世界モデルとは、エージェントが内部に持つ、世界の仕組みに関する抽象的かつ予測的な表現である。これを持つことで、エージェントは行動の結果をシミュレートし、複雑な計画を立て、原因と結果について推論することが可能になる。これらは、現在の LLM が決定的に欠いている能力である³⁸。ルカン氏は、「我々は猫の知能すら再現できていないのに、人間の知能を語ることはできない」と述べ、テキストデータのみから学習するモデルの限界を鋭く指摘している³⁸。

技術アーキテクチャ：JEPA

世界モデル構築に向けた具体的な一歩として、ルカン氏が提案するのが「共同埋め込み予測アーキテクチャ（Joint-Embedding Predictive Architecture, JEPA）」である³⁹。これは、ピクセルやトークンといった生の入力空間で予測を行う従来の生成モデルとは一線を画す。JEPA は、入力データ（例えば画像の一部）から高次元の抽象的な表現（埋め込み）を生成し、その埋め込み空間内で、他の部分の表現を予測するように学習する⁴¹。このアプローチにより、モデルは表面的な詳細（例：草の葉一枚一枚）を再構築するのではなく、より本質的で抽象的な概念（例：「野原」という概念）を学習することが期待される。これは、効率的な世界モデルを構築するための重要なステップと考えられている。

身体性（エンボディメント）のテーゼ

世界モデルの概念は、「身体性（Embodiment）」のテーゼと密接に結びついている。このテーゼの核心は、頑健な世界モデルはテキストデータのような受動的な情報からだけでは学習できず、物理世界との能動的な相互作用を通じて「接地（グラウンディング）」される必要があるという主張である⁴⁴。エージェントが物理法則、因果関係、空間認識といった世界の根源的なルールを真に理解するためには、センサーやアクチュエーターを備えた身体を持ち、試行錯誤を通じて世界を経験する必要がある⁹。この観点から、身体性を持つ AI（Embodied AI）の研究では、その能力に応じてレベル 1（単一タスク遂行）からレベル 5（全目的ロボット）までの段階的なロードマップが提案されている⁴⁶。

4.2 ニューロシンボリック統合

中核概念：ゲイリー・マーカスの批判

認知科学者であり、AI に対する鋭い批評家として知られるゲイリー・マーカス氏は、AI が人間のような知能を獲得するためには、歴史的に対立してきた 2 つのアプローチを統合する必要があると主張する。それが「ニューロシンボリック AI (Neuro-Symbolic AI)」である³¹。このアプローチは、データからの学習やパターン認識に長けたニューラルネットワーク（ダニエル・カーネマンの言う「システム 1」に相当）と、論理、ルール、記号操作に基づく厳密な推論を得意とするシンボリック AI（「システム 2」に相当）を組み合わせることを目指す。

LLM の欠陥への対処

このハイブリッドなアプローチは、LLM が抱える多くの弱点に直接的に対処する可能性を秘めている。シンボリックなコンポーネントを導入することで、AI システムに構造化された知識ベース（ナレッジグラフなど）を提供し、論理的な一貫性を強制し、そして AI の推論プロセスを人間が理解できる形で透明化・解釈可能にすることができる³⁰。これにより、ハルシネーションを抑制し、信頼性の高い推論を実現することが期待される。

応用とアーキテクチャ

ニューロシンボリック AI の具体的な応用例はすでに見られる。例えば、自動運転車は、ニューラルネットワークを使ってカメラ映像から物体（歩行者、他の車、信号機）を認識し（ニューラル部分）、同時にシンボリックなエンジンを使って交通法規というルールに基づいて行動を決定する（シンボリック部分）⁵³。また、Google DeepMind が開発した「AlphaGeometry」は、LLM の直感的な発想力と、シンボリックな演繹エンジンによる厳密な証明能力を組み合わせることで、国際数学オリンピックレベルの幾何学の問題を解くことに成功した⁵³。アーキテクチャとしては、シンボリックな入力をニューラルネットワークで処理し、再びシンボリック

な出力に戻す「Symbolic → Neuro → Symbolic」のような逐次的なモデルなどが研究されている⁵¹。

4.3 意識の難問

中核概念：ヨシュア・ベンジオのアプローチ

AI 研究のもう一人の重鎮であるヨシュア・ベンジオ氏は、AGI のパズルを解く上で、意識の理解が重要な役割を果たす可能性があるという視点を提供している。彼のアプローチは、意識を神秘的なものとして扱うのではなく、神経科学の知見に基づき、意識の機能を計算論的に解明し、その「指標となる特性 (indicator properties)」を特定しようとする、厳密で科学的なものである⁵⁴。

機能主義 vs. 生物学的自然主義

この探求は、意識に関する哲学的な大論争に直結する。一方には、意識は情報処理の「機能」によって定義され、その機能を実装すれば、シリコンのような非生物学的な基盤の上でも再現可能であるとする「計算論的機能主義 (computational functionalism)」がある⁵⁶。もう一方には、意識は脳の神経細胞や生化学的プロセスといった、生命システム特有の物理的・生物学的特性から生じる固有の現象であるとする「生物学的自然主義 (biological naturalism)」がある⁵⁶。

AGI への関連性

この議論は、単なる哲学的な思弁にとどまらない。もし、人間が持つ高度な推論、汎化、あるいは道徳的判断能力の一部が、意識を持つことと不可分であるならば、「真の」AGI を構築するためには、意識を持つ AI を作ることが前提条件となるかもしれない。さらに、たとえ技術的に意識の再現が不可能であったとしても、人間が AI を「意識があるかのように認識する」だけ

で、社会に深刻な倫理的・法的な混乱を引き起こす可能性がある⁵⁵。したがって、意識の問題は、AGI の技術的実現可能性と社会的受容性の両方にとって、避けては通れない課題なのである。

これらの代替パラダイムは、一見すると互いに競合しているように見えるが、より深いレベルでは相互補完的であり、収斂する可能性を秘めている。ルカン氏の世界モデルは、マーカス氏が提唱するシンボリックな構造によって、その知識表現がより豊かで頑健になるかもしれない。そして、その世界モデルは、身体性を伴うエージェントが実世界で経験を積むことによって初めて、真に意味のある形で接地されるだろう。このことから、AGI への最も有望な道筋は、これら 3 つの要素——身体性を持ち、ニューロシンボリックな世界モデルを構築するエージェント——の統合にあるのかもしれない。各陣営は、それぞれが巨大で複雑なパズルの、必要不可欠なピースを記述しているのである。

第 5 章 哲学的な行き詰まり：機械は真に理解できるか？

AGI をめぐる技術的な議論の根底には、より深く、解決が困難な哲学的な問いが存在する。それは、「機械は、人間がするように、真に意味を『理解』することができるのか？」という問いである。この問いは、AI の能力の限界を定めるだけでなく、「知能」そのものの本質に迫るものであり、ジョン・サールの「中国語の部屋」の思考実験によって鮮やかに提起された。

「中国語の部屋」の再訪

1980 年に発表されたサールの思考実験は、次のようなシナリオを提示する³。部屋の中に、中国語を全く理解できない英語話者の男性が一人いる。彼の前には、英語で書かれた膨大なルールの本（プログラム）と、たくさんの中国語の記号（データベース）がある。部屋の外から、中国語で書かれた質問がスリットを通して差し入れられる。男性は、ルールの本に従って、受け取った記号に対応する別の記号を探し出し、それを外に返す。部屋の外にいる中国語話者から見れば、部屋はまるで中国語を完全に理解しているかのように、質問に対して適切で流暢な回答を返しているように見える。

しかし、サールは問う。この部屋の中に「理解」は存在するだろうか？ 男性は記号を操作しているだけで、その意味を全く理解していない。ルールの本も、記号の箱も、それ自体が何かを理解しているわけではない。サールの結論は、このシステム全体（男性＋本＋記号）がいかに知的な振る舞いを見せようとも、そこには真の「理解」や「志向性（intentionality）」は存在

しない、というものである。彼は、コンピュータが行っていることは、本質的にこの部屋の中の男性が行っている記号操作（統語論、シンタックス）と同じであり、それだけでは意味の理解（意味論、セマンティクス）は決して生まれないと主張した⁶²。

LLM の時代における議論の妥当性

この思考実験は、現代の LLM に直接的に適用できる。LLM は、その巨大な規模と複雑さにもかかわらず、本質的には非常に高度な「中国語の部屋」であると見なすことができる⁶¹。LLM は、膨大な学習データの中から統計的なパターン（統語論的なルール）を抽出し、与えられた入力（プロンプト）に対して、最も確率の高い記号（トークン）の系列を生成している。しかし、それらの記号が現実世界で何を指し示しているのか（意味論）を、身体的な経験を通じて接地（グラウンディング）しているわけではない。これは認知科学で「シンボル・グラウンディング問題」として知られる課題であり、LLM が直面する根源的な問題である⁶³。

根深い対立：機能主義 vs. 生物学的自然主義

「中国語の部屋」が浮き彫りにした対立は、心の哲学における二つの主要な立場、すなわち「機能主義」と「生物学的自然主義」の間の根深い溝を反映している。

- **機能主義（Functionalism）**：この立場は、心的な状態（思考、信念、理解など）は、その物理的な構成要素（脳のニューロンか、コンピュータのトランジスタか）によってではなく、システム内で果たす「機能」や「因果的役割」によって定義されると主張する³。機能主義者にとって、もしシリコンベースのシステムが、人間の脳と全く同じ情報処理機能を実現できるのであれば、そのシステムは心を持つと言える。この哲学は、「強い AI」や、現在のスケーリング仮説の支持者たちが暗黙のうちに前提としている考え方である。
- **生物学的自然主義（Biological Naturalism）**：これはサール自身の立場であり、意識や理解といった心的現象は、脳のニューロンや生化学的プロセスが持つ特定の生物学的な因果力から生じる、具体的な生物学的現象であると主張する⁵⁶。この見解によれば、心をシミュレーションすることと、心を持つことは全く異なる。コンピュータ上で嵐をシミュレーションしても部屋が濡れることはないように、脳の働きをシミュレーションしても、それだけでは意識や理解は生まれない。

この哲学的な対立は、AGI をめぐる技術的な議論が単なる手法の優劣を競うものではないことを示している。それは、「知能とは何か」という本質的な問いに対する、根本的な世界観の対

立なのである。一方の陣営は、知能を基盤（サブストレート）に依存しない純粋な情報処理、すなわち一種の計算であると捉えている。この立場に立てば、AGI の実現は、十分な計算能力と優れたアルゴリズムを開発する工学的な課題となる。

もう一方の陣営は、知能、特に意識や理解は、環境と相互作用する生きた身体を持つ有機体（オーガニズム）の創発的な特性であると主張する。この立場に立てば、AGI の実現は、単に情報処理装置を構築するだけでなく、ある種の人工的な生命体を創造する課題に近づいていく。

もちろん、サールの議論には多くの反論が寄せられてきた。例えば、「システム応答（The Systems Reply）」は、部屋の中の男性は中国語を理解していないが、男性、ルールの本、記号の箱からなる「システム全体」としては中国語を理解していると主張する³。また、LLM の学習能力や汎化能力は、静的なルールブックに従うだけの男性とは根本的に異なるとも指摘されている⁶⁵。そして前述の通り、ジェフリー・ヒントンは、十分な予測能力は必然的に理解を伴うとし、主観的経験と客観的能力の区別自体を否定している²⁷。

しかし、これらの反論があってもなお、「中国語の部屋」が提起した問題の核心は残る。現在の AI が示す知的な振る舞いと、我々が人間に対して用いる「理解」という言葉との間には、依然として埋めがたい質的な断絶があるのではないか。この哲学的な行き詰まりは、AGI への道が単一の技術的ブレークスルーによって開かれるのではなく、知能、意識、そして生命そのものについての我々の理解が深まることと並行して、ゆっくりと進んでいくであろうことを示唆している。

第 6 章 統合と戦略的展望：収斂する未来へ

本レポートで詳述してきたように、「スケーリングのみで AGI は達成可能か」という問いは、技術的、理論的、そして哲学的な側面が複雑に絡み合った、現代における最重要課題の一つである。分析の結果、この問いを「はい/いいえ」の二元論で捉えることは、事態の本質を見誤らせる可能性があることが明らかになった。未来は、どちらか一方のアプローチが勝利するのではなく、複数のパラダイムが収斂し、統合される方向へと進む可能性が極めて高い。

二元論の超克：統合的アプローチの必然性

「スケーリング vs. 新アーキテクチャ」という対立構造は、議論を明確にする上では有用だ

が、現実の AI 開発の未来を予測する上では誤解を招きかねない。最も可能性の高いシナリオは、両者が互いの長所を補い合う形で統合される未来である。スケールアップされた LLM が持つ、広範な知識が圧縮された巨大な基盤は、強力な「システム 1」（直感的・連想的な思考）として機能するだろう。そして、その上に、世界モデルやニューロシンボリック・エンジンといった、より構造化され、論理的な推論を司る「システム 2」が構築される⁵⁰。この統合により、LLM の柔軟性と知識の広さと、シンボリックシステムの厳密性と信頼性を兼ね備えた、より堅牢で汎用的な知能が生まれる可能性がある。

AGI への予測される軌道

この統合的アプローチを念頭に置くと、AGI への道筋は以下のような多段階の進化をたどると予測される。

1. **短期（1～3 年）**：スケーリングは継続するが、その主眼は純粋な規模の拡大から、学習効率やエネルギー効率の向上へとシフトする。「学習後改善」技術（強化学習、推論時計算など）への投資がさらに加速し、モデルの信頼性と能力の向上が図られる¹⁷。複数のステップからなるタスクを自律的に実行できる、より高度な「エージェント的 AI」が登場し、特定の業務領域で実用化が進むだろう²⁶。
2. **中期（3～10 年）**：ハイブリッドモデルが本格的に商用化される時代。LLM をナレッジグラフやシンボリックな推論エンジンと明示的に組み合わせたシステムが、エンタープライズ分野で広く採用されるようになる。物理世界で活動する身体性 AI（ロボティクス）は、研究室レベルから、物流や製造といったニッチな商業応用へと進出する⁴⁶。この段階では、異なるアーキテクチャ間の「接着剤」としての技術が重要となる。
3. **長期（10 年以上）**：身体的な相互作用を通じて、接地された世界モデルを自己学習する、真に新しい統合アーキテクチャの探求が本格化する。これは、より堅牢で、真に汎用的な知能への道筋であるが、その実現時期は極めて不確実である²⁹。この段階でのブレークスルーは、現在の AI 技術とは質的に異なる、根本的な科学的発見を必要とするかもしれない。

主要な AGI 哲学の比較分析

この複雑な議論の全体像を把握するために、本分野を牽引する主要な思想家たちの立場を以下の表にまとめる。この表は、彼らのアプローチ、信念、そして懸念事項の間の類似点と相違点を浮き彫りにし、AGI をめぐる思想的ランドスケープを明確に描き出す。

次元	サム・アルトマン (OpenAI)	ジェフリー・ヒント ン	ヤン・ルカ ン (Meta)	ゲイリー・ マーカス	ヨシュア・ ベンジオ
AGI への主 要経路	反復的スケ ーリング： モデル、デ ータ、計算 量の継続的 なスケール アップが主 要な駆動 力。AGI は 継続的な改 善プロセス である。	スケーリン グ＋創発： スケーリン グは真の創 発的理解に つながる。 基本原理は 健全であ る。	新アーキテ クチャ：現 行の LLM は行き止ま り。世界モ デルと、生 成モデルか らの脱却が 必要。	ハイブリッ ドシステ ム：ニュー ラルネット ワークと古 典的なシン ボリック AI の統合が必 要。	神経科学に 着想を得た AI：意識と 高次認知の 計算論的原 理の理解が 必要。
LLM の 「理解」に 対する見解	実用的／暗 黙の肯定： 複雑なタス クをこな し、科学を 加速させる モデルは、 ある種の理 解を示して いる。	明確な肯 定：大規模 な次単語予 測能力は、 真の理解を 必然的に伴 う。主観と 客観の区別 は「くだら ない」。	明確な否 定：LLM は世界モデ ルを欠いて おり、推 論、計画、 真の理解は 不可能。表 面的な統計 に限定され る。	明確な否 定：LLM は模倣に長 けた「確率 的オウム」 であり、因 果的推論、 世界モデ ル、信頼性 を欠く。	慎重／懐疑 的：現行シ ステムは意 識を持た ず、脳の意 識的处理に 関連する主 要な計算特 性を欠いて いる。
提案される 次の一手／ アーキテク チャ	さらなるス ケール (GPT- n)：計算 量とデータ への大規模 投資。RL	アナログ・ コンピュー ティング： 知能のスケ ールを継続 するため、 より脳に近	JEPA （共 同埋め込み 予測アーキ テクチャ）： 抽象 的な世界モ デルを学習	ニューロシ ンボリッ ク・ハイブ リッド：ニ ューラルネ ットをシン ボリックな	接地された AI 研究： 神経科学の 理論から導 出された、 意識の「指 標となる特

	やエージェンツのチューニングによる学習後改善を組み合わせる。	く、エネルギー効率の良いハードウェアを探索する。	する非生成モデル。身体性を通じて接地される可能性が高い。	推論エンジン、ナレッジグラフ、論理エンジンと統合する。	性」を持つシステムを開発する。
主要な懸念／焦点	利益の分配と社会の適応： AGI が全人類に利益をもたらすことを保証し、社会経済的な移行を管理すること。	実存的リスクと制御： AI が人間より賢くなり、意図しない権力志向のサブゴールを追求する可能性。	常識の欠如： 動物にできる基本的な推論や計画が、現在の AI にはできないこと。	信頼性と信用性： LLM に内在する欠陥（幻覚、バイアス）が、高リスクな応用を不可能にしていること。	倫理的・社会的混乱： 人間が AI を意識があると「認識」することのリスクと、厳格な安全性・アライメント研究の必要性。

結論

AGI の探求は、21 世紀における最も重大な技術的・哲学的事業の一つである¹³。その道のりは、目的地そのものと同じくらい重要である。なぜなら、それは我々自身が持つ知能、意識、そして価値観といった根源的な問いと向き合うことを強いるからだ⁶⁹。

本レポートの分析が示すように、AGI への道は単線的ではない。それは、スケーリングという力強い潮流と、世界モデル、ニューロシンボリック、身体性といった新しいアーキテクチャへの分岐が交差する、複雑な地形をしています。前進するためには、技術的な才能だけでなく、分野を超えた協力、倫理的な先見性、そしてこの多面的な挑戦に対する深い理解が不可欠である⁷²。未来は、一つのアプローチが他を圧倒するのではなく、異なるアイデアの賢明な統合によって築かれるだろう。その統合をいかにして成し遂げるかが、AGI の夢を実現するための鍵となる。

引用文献

1. 強い AI と AGI : なぜ意識は定義から外れたのか | yacc - note, 10 月 12, 2025 にアクセス、<https://note.com/cutaway/n/nc46761fc1b86>
2. 汎用人工知能 (AGI) | 用語解説 | 野村総合研究所(NRI), 10 月 12, 2025 にアクセス、<https://www.nri.com/jp/knowledge/glossary/agi.html>
3. Chinese room - Wikipedia, 10 月 12, 2025 にアクセス、https://en.wikipedia.org/wiki/Chinese_room
4. 第 1 回 : AGI の定義の変遷と理論的議論 - note, 10 月 12, 2025 にアクセス、https://note.com/reini_mizushima/n/nfb76406ad17d
5. 汎用人工知能 - Wikipedia, 10 月 12, 2025 にアクセス、<https://ja.wikipedia.org/wiki/%E6%B1%8E%E7%94%A8%E4%BA%BA%E5%B7%A5%E7%9F%A5%E8%83%BD>
6. 汎用人工知能 (AGI) 完全ガイド : 他の AI との違い・できること、社会的悪影響、開発事例と準備しておくべきこと - Relipa, 10 月 12, 2025 にアクセス、https://relipasoft.com/blog/complete_guide-to-artificial-general-intelligence-agi/
7. AGI とは? 人間を超える知性を変える未来と私たちの選択肢を解説 - 株式会社アドカル, 10 月 12, 2025 にアクセス、<https://www.adcal-inc.com/column/agi-introduction/>
8. AGI(汎用人工知能)とは? AI や ChatGPT との関係性・社会的課題 ..., 10 月 12, 2025 にアクセス、https://www.brainpad.co.jp/doors/contents/about_agi/
9. AGI (汎用人工知能) とは? AI・ASI との違いや特徴・活用例を紹介 | OPTAGE for Business, 10 月 12, 2025 にアクセス、https://optage.co.jp/business/contents/article/what_is-agi-asi.html
10. Levels of AGI for Operationalizing Progress on the Path to AGI - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/pdf/2311.02462>
11. AGI and consciousness: are we safe?- MedCrave online, 10 月 12, 2025 にアクセス、https://medcraveonline.com/AHOAJ/agi_and-consciousness-are-we-safenbsp.html
12. AGI 実現までのロードマップを歩み出す | AI 専門ニュースメディア AINOW, 10 月 12, 2025 にアクセス、<https://ainow.ai/2024/02/02/275671/>
13. AGI とは? 人類の知能を超えるかもしれない次世代 AI の全貌と社会へのインパクトを徹底解説, 10 月 12, 2025 にアクセス、<https://www.macromill.com/service/words/agi/>
14. medium.com, 10 月 12, 2025 にアクセス、<https://medium.com/bossier-tech/the-scaling-law-formula-80209fab191d#:~:text=The%20Fundamental%20Idea%20Behind%20Scaling%20Laws&text=At%20its%20core%2C%20the%20scaling,biology%2C%20and%20even%20human%20learning.>
15. Neural Scaling Laws For AGI : r/learnmachinelearning- Reddit, 10 月 12, 2025 にアクセス、https://www.reddit.com/r/learnmachinelearning/comments/1dy1ldz/neural_scaling

[g laws for agi/](#)

16. Scaling Laws for Neural Language Models - arXiv, 10 月 12, 2025 にアクセス、
<https://arxiv.org/pdf/2001.08361>
17. The Future of LLMs and AGI- AIM Councils, 10 月 12, 2025 にアクセス、
<https://councils.aimediahouse.com/the-future-of-llms-and-agi/>
18. The Race to Efficiency: A New Perspective on AI Scaling Laws - arXiv, 10 月 12, 2025 にアクセス、
<https://arxiv.org/html/2501.02156v3>
19. [2503.05788] Emergent Abilities in Large Language Models: A Survey - arXiv, 10 月 12, 2025 にアクセス、
<https://arxiv.org/abs/2503.05788>
20. Emergent Abilities in Large Language Models: A Survey - arXiv, 10 月 12, 2025 にアクセス、
<https://arxiv.org/html/2503.05788v1>
21. Emergent Abilities in Large Language Models: A Survey - arXiv, 10 月 12, 2025 にアクセス、
<https://arxiv.org/pdf/2503.05788>
22. Why are LLMs' abilities emergent? - arXiv, 10 月 12, 2025 にアクセス、
<https://www.arxiv.org/pdf/2508.04401>
23. What Does Hitting Scaling Law Limit Mean for US-China AI Competition - Interconnected, 10 月 12, 2025 にアクセス、
<https://interconnected.blog/what-does-hitting-scaling-law-limit-mean-for-us-china-ai-competition/>
24. Three Observations - Sam Altman, 10 月 12, 2025 にアクセス、
<https://blog.samaltman.com/three-observations>
25. Sam Altman says GPT-5 critics are wrong, calls it a big leap toward real scientific AI, 10 月 12, 2025 にアクセス、
<https://www.indiatoday.in/technology/news/story/sam-altman-says-gpt-5-critics-are-wrong-calls-it-a-big-leap-toward-real-scientific-ai-2798430-2025-10-06>
26. Reflections - Sam Altman, 10 月 12, 2025 にアクセス、
<https://blog.samaltman.com/reflections>
27. Can AGI Think? (Geoff Hinton) – Rob Schläff's Website, 10 月 12, 2025 にアクセス、
<https://schlaaff.com/wp/almanac/things-i-like/technical-ideas/can-agi-think-geoff-hinton/>
28. What does Geoffrey Hinton believe about AGI existential risk? - Marginal REVOLUTION, 10 月 12, 2025 にアクセス、
<https://marginalrevolution.com/marginalrevolution/2023/06/what-does-geoffrey-hinton-believe-about-agi-existential-risk.html>
29. AI や AGI の急速な発展に対する慎重な意見や懸念について、著名な専門家たちの見解をまとめます, 10 月 12, 2025 にアクセス、
<https://www.acrovision.jp/suemitsu/?p=916>
30. Generative AI's crippling and widespread failure to induce robust models of the world, 10 月 12, 2025 にアクセス、
<https://garymarcus.substack.com/p/generative-ais-crippling-and-widespread>
31. 'Not on the Best Path' – Communications of the ACM, 10 月 12, 2025 にアクセス、
<https://cacm.acm.org/opinion/not-on-the-best-path/>

32. 'World Models,' an Old Idea in AI, Mount a Comeback | Quanta ..., 10 月 12, 2025 にアクセス、<https://www.quantamagazine.org/world-models-an-old-idea-in-ai-mount-a-comeback-20250902/>
33. A Comparative Study of Neurosymbolic AI Approaches to Interpretable Logical Reasoning, 10 月 12, 2025 にアクセス、<https://arxiv.org/html/2508.03366v1>
34. 24%が「AGI が来た」と思えば、それが AGI。後悔も定義も、すべて人間の解釈次第。 - サートプロ, 10 月 12, 2025 にアクセス、<https://www.certpro.jp/blogs/20250404-1/>
35. AI 研究者は AGI の実現に懐疑的？AI 学会の最新調査 - XenoSpectrum, 10 月 12, 2025 にアクセス、<https://xenospectrum.com/is-it-difficult-to-achieve-agi-just-by-scaling-up-current-ai/>
36. Scaling Laws: It's GPT-5 Against the Human Brain, not AI/AGI - WorldHealth.net, 10 月 12, 2025 にアクセス、<https://worldhealth.net/news/gpt-5-against-the-human-brain-not-ai-agi/>
37. 【シンギュラリティを問う Vol.3】 AI 進化の壁と可能性。シンギュラリティと私たちの未来 - Salesforce, 10 月 12, 2025 にアクセス、<https://www.salesforce.com/jp/blog/jp-hakuhodo-mori-singularity-vol3/>
38. Tech leaders eye world models as link to smarter AI - IBM, 10 月 12, 2025 にアクセス、<https://www.ibm.com/think/news/world-models-smarter-ai>
39. Critical review of LeCun's Introductory JEPa paper | Medium - Malcolm Lett, 10 月 12, 2025 にアクセス、<https://malcolmlett.medium.com/critical-review-of-lecuns-introductory-jepa-paper-fabe5783134e>
40. A Path Towards Autonomous Machine Intelligence Version 0.9.2, 2022-06-27 - OpenReview, 10 月 12, 2025 にアクセス、<https://openreview.net/pdf?id=BZ5a1r-kVsf>
41. LLM-JEPA: Large Language Models Meet Joint Embedding Predictive Architectures - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/abs/2509.14252>
42. LLM-JEPA: Large Language Models Meet Joint Embedding Predictive Architectures - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/html/2509.14252v2>
43. Could someone explain what each of these architectures are that LeCun claims could lead to AGI? : r/singularity - Reddit, 10 月 12, 2025 にアクセス、https://www.reddit.com/r/singularity/comments/likpqyk/could_someone_explain_what_each_of_these/
44. 通用人工知能(AGI) 事例 - IBM, 10 月 12, 2025 にアクセス、<https://www.ibm.com/cn-zh/think/topics/artificial-general-intelligence-examples>
45. AGI(汎用人工知能) とは何ですか? - AWS, 10 月 12, 2025 にアクセス、<https://aws.amazon.com/jp/what-is/artificial-general-intelligence/>
46. [Literature Review] Toward Embodied AGI: A Review of Embodied AI and the Road Ahead, 10 月 12, 2025 にアクセス、<https://www.themoonlight.io/en/review/toward-embodied-agi-a-review-of->

[embodied-ai-and-the-road-ahead](#)

47. Is Embodied Humanoid AI a Tipping Point Toward AGI? | by ..., 10 月 12, 2025 にアクセス、<https://medium.com/@innovationstrategy/embodied-ai-thinking-machines-c16066dbf8df>
48. Embodied Intelligence: The Key to Unblocking Generalized Artificial Intelligence - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/html/2505.06897v1>
49. Toward Embodied AGI: A Review of Embodied AI and the Road Ahead - Reading Notes, 10 月 12, 2025 にアクセス、<https://lsy641.github.io/notes/towards-embodied-AI>
50. 汎知能 (AGI) とは? | ガートナー ジャパン (Gartner) , 10 月 12, 2025 にアクセス、<https://www.gartner.co.jp/ja/topics/emerging-technology-watch-agi>
51. Unlocking the Potential of Generative AI through Neuro-Symbolic Architectures – Benefits and Limitations - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/html/2502.11269v1>
52. Building Trustworthy NeuroSymbolic AI Systems: Consistency, Reliability, Explainability, and Safety - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/html/2312.06798v1>
53. Decoding Neuro-Symbolic AI - Phaneendra Kumar Namala - Medium, 10 月 12, 2025 にアクセス、<https://phaneendrkn.medium.com/decoding-neuro-symbolic-ai-64385310f030>
54. Consciousness in Artificial Intelligence: Insights from the Science of Consciousness - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/abs/2308.08708>
55. SCIENCE. Illusions of AI consciousness. The belief that AI is conscious is not without risk. Yoshua Bengio and Eric Elmoznino - blog.biocomm.ai, 10 月 12, 2025 にアクセス、<https://blog.biocomm.ai/2025/09/14/science-illusions-of-ai-consciousness-the-belief-that-ai-is-conscious-is-not-without-risk-yoshua-bengio-and-eric-elmoznino/>
56. Artificial Intelligence as an Opportunity for the Science of Consciousness - arXiv, 10 月 12, 2025 にアクセス、<https://arxiv.org/pdf/2509.07001>
57. The Functionalist Case for Machine Consciousness: Evidence from Large Language Models, 10 月 12, 2025 にアクセス、<https://www.lesswrong.com/posts/Hz7igWbjS9joYjfDd/the-functionalist-case-for-machine-consciousness-evidence>
58. Artificial consciousness and biological naturalism: a perspective between computation, living dynamics, and ethical considerations - Andrea Viliotti, 10 月 12, 2025 にアクセス、<https://www.andreaviliotti.it/post/artificial-consciousness-and-biological-naturalism-a-perspective-between-computation-living-dynamics>
59. AGI を深く理解する | 朱雀 | SUZACQUE - note, 10 月 12, 2025 にアクセス、<https://note.com/suzacque/n/n7935c4e07964>
60. The consciousness question in the age of AI - Cosmos Magazine, 10 月 12, 2025 にアクセス、<https://cosmosmagazine.com/my-cosmos/ai-consciousness-my>

[cosmos/](#)

61. How Searle's Chinese Room Contends with Modern LLMs, 10 月 12, 2025 にアクセス、
<https://thopay.dev/blog/how-searles-chinese-room-contents-with-modern-llms>
62. The Chinese Room Argument - Stanford Encyclopedia of Philosophy, 10 月 12, 2025 にアクセス、
<https://plato.stanford.edu/entries/chinese-room/>
63. Is the "Chinese room" an explanation of how ChatGPT works? - AI Stack Exchange, 10 月 12, 2025 にアクセス、
<https://ai.stackexchange.com/questions/39293/is-the-chinese-room-an-explanation-of-how-chatgpt-works>
64. Conscious artificial intelligence and biological naturalism - PubMed, 10 月 12, 2025 にアクセス、
<https://pubmed.ncbi.nlm.nih.gov/40257177/>
65. Generative AI vs The Chinese Room Argument : r/singularity - Reddit, 10 月 12, 2025 にアクセス、
<https://www.reddit.com/r/singularity/comments/17te9yn/generative-ai-vs-the-chinese-room-argument/>
66. 大規模言語モデルは AGI に繋がるか？ | 山川宏が語る全脳アーキテクチャの進化と AI の未来, 10 月 12, 2025 にアクセス、
<https://mirai-agenda.com/posts/yanakawa1>
67. Sam Altman Predicts AGI by 2027, AI Agents to Reshape Workforces - WebProNews, 10 月 12, 2025 にアクセス、
<https://www.webpronews.com/sam-altman-predicts-agi-by-2027-ai-agents-to-reshape-workforces/>
68. なぜ、AGI の登場は今から 30 年後くらいになるのか？ - Zenn, 10 月 12, 2025 にアクセス、
<https://zenn.dev/pdfractal/articles/0fa85a99ab1152>
69. AGI/シンギュラリティ論争は「無意味」？建設的な AI 議論への道筋 | Shicci-note, 10 月 12, 2025 にアクセス、
<https://note.com/shicci/n/nef69fa678d2c>
70. 人間学習の新地平、汎用 AI 時代に再定義される教育 - AuthenticAI, 10 月 12, 2025 にアクセス、
<https://authentica.co.jp/blogs/contents/a-new-horizon-in-human-learning-education-redefined-in-the-era-of-general-purpose-ai>
71. AI 全盛時代における哲学的考察や洞察の価値 | インディ・パ | 生成 AI 教育・研修・コンサルティング, 10 月 12, 2025 にアクセス、
<https://indepa.net/archives/7447>
72. AGI とは？従来の AI との違いや活用領域、今後の課題を徹底解説 | AI 総合研究所, 10 月 12, 2025 にアクセス、
<https://www.ai-souken.com/article/what-is-agi>