



このままAIが発展するだけでAGI達成可能か：スケーリング仮説をめぐる包括的分析

提示されたYouTube動画「このままAIが発展するだけでAGI達成可能な理由」とスケーリング仮説に関する国際的な議論を深掘りした結果、AI研究コミュニティは現在、人工汎用知能（AGI）の実現経路について根本的に異なる2つの陣営に分かれていることが判明しました。本分析では、スケーリング仮説の根拠と限界、代替アプローチ、そして議論の根底にある哲学的前提について包括的に検証します。

動画・資料分析：スケーリング仮説の核心的論理

提示されたYouTube動画では、Frieveが「スケーリング則」（Scaling Laws）に基づき、現在のAI技術の延長線上でAGIが実現可能であると主張しています。動画の核心的論拠は以下の通りです：^[1]

スケーリング則の数学的根拠：動画では、Kaplan et al.（2020年）の研究を引用し、AIモデルの誤差（L）が資源（N）に対して $L = (N_c/N)^\alpha$ の冪法則に従うことを説明しています。この法則によれば、モデルサイズ、データ量、計算量を増やすことで、予測可能な形で性能が向上し続けます。^[2]^[3]

人間の学習との対比：動画では、人間の学習能力には生物学的制約（80年の寿命、1日10時間程度の学習時間、脳容量の限界）があるのに対し、AIは理論的に無制限にこれらの資源を拡張できると論じています。これにより、人間を超える性能の実現が可能になるとしています。

効率化によるスケーリング継続：2025年に入り、純粋な計算量増加だけでなく、効率化技術の進歩により、予想よりも少ない資源で同等の性能を実現できるようになったことを「石油発掘技術の進歩」に例えて説明しています。

スケーリング仮説支持派の論拠と証拠

現在のAIアプローチの延長でAGI達成可能とする肯定的見解は、複数の実証的証拠に基づいています：^[4] ^[5] ^[6]

実証的スケーリング則：OpenAIのKaplanらが2020年に発表した研究では、言語モデルの損失がモデルサイズ、データセット、計算量に対して7桁以上の範囲で一貫した冪法則を示すことが実証されました。具体的には、パラメータ数に対して $\alpha_N \sim 0.076$ 、データサイズに対して $\alpha_D \sim 0.095$ の指数で性能が向上します。^[3] ^[2]

テスト時計算スケーリング：2024年以降注目されているのが、推論時の計算量を増やすことで性能が向上する現象です。OpenAIのo1モデルのように、回答を生成する際により多くの計算時間をかけることで、より高い精度を実現できることが示されています。^[7]

マルチモーダル・スケーリング：単なるテキストデータの増加だけでなく、画像、音声、動画など複数のモダリティを統合することで、さらなる性能向上が期待されています。これにより、人間のよう

な多感覚統合能力の実現可能性が示唆されています。^[7]

エマージェント能力の発現：大規模言語モデルの開発過程で、訓練時には明示的に学習していない能力（数学的推論、プログラミング、創作など）が一定の規模を超えると突然現れる現象が観察されています。^[8]

スケーリング仮説への主要な反論と批判

一方で、スケーリング仮説に対する批判は多角的かつ根本的です：^{[9] [10] [11]}

構造的理解の欠如：カリフォルニア大学バークレー校のStuart Russellは、現在のAIアーキテクチャが「大規模なフィードフォワード回路」に基づいており、概念を表現する能力に根本的な制限があると指摘しています。これらのシステムは断片的な表現しか生成できず、真の理解には至らないとされています。^[9]

3つの壁理論：スケーリングの限界として、以下の「3つの壁」が指摘されています：^[11]

- **データの壁**：高品質な人間作成データの枯渇
- **計算・エネルギーの壁**：物理的・経済的制約
- **アーキテクチャの壁**：次トークン予測の根本的限界

研究者コンセンサスの変化：2025年のアメリカ人工知能学会による調査では、475人のAI研究者のうち76%が「現在の大規模言語モデルをスケールアップしてもAGIを開発できる可能性は低い」と回答しています。これは、専門家コミュニティ内でスケーリング仮説への懐疑論が主流になっていることを示しています。^[9]

議論の主要人物と研究機関の立場

スケーリング仮説支持派

Sam Altman (OpenAI CEO)：「深層学習が機能し、スケールと共に予測可能に改善され、我々はそれに増大するリソースを投入した」として、スケーリングの有効性を強調しています。ただし、最近では「LLMだけでは不十分で、さらなるブレークスルーが必要」とも発言しており、純粋なスケーリング論から距離を置き始めています。^{[12] [13]}

Dario Amodei (Anthropic CEO)：スケーリング仮説を「料理の材料を増やすほど豊かな味わいになる」と例え、より多くのデータと処理能力を使用することで、モデルがより賢く、より有能になると主張しています。^[5]

Gwern Branwen：「スケーリング仮説」の主要な理論家として、ニューラルネットワークがデータと計算を吸収し、問題が難しくなるにつれてより一般化しベイズ的になり、新しい能力を発現するという理論を提唱しています。^[14]

スケーリング仮説懐疑派

Yann LeCun (Meta Chief AI Scientist)：現在のLLMベースのアプローチに最も批判的な研究者の一人です。彼は「大規模言語モデルは人間レベルの知能に到達しない。AGI達成は技術開発問題ではなく、非常に科学的な問題だ」と述べています。LeCunはJEPA (Joint Embedding Predictive

Architecture) と世界モデルを提唱し、言語中心から物理世界理解中心へのパラダイムシフトを主張しています。[15] [16] [17] [18]

Gary Marcus (NYU名誉教授)：長年にわたってディープラーニングとスケーリング仮説の限界を指摘してきました。彼は「スケーリング則は普遍的法則ではなく、必ずしも成り立たない観察結果に過ぎない」と主張し、真の知能には記号的推論と常識的理解が不可欠であると論じています。[19] [20] [21]

Stuart Russell (UC Berkeley)：AIの安全性研究の第一人者として、現在のAIシステムの不透明性と制御の困難さを強調しています。彼は「現在のAIシステムの内部動作原理は不明であり、安全性と規制に深刻な問題をもたらす」と警告しています。[22] [23]

Song-Chun Zhu (BIGAI所長)：中国の清華大学でAGI研究を主導し、「我々は社会として『AI』という用語を誤解している可能性がある。多機能携帯電話を『スマート』と呼ぶように、今日使われている人気のAIモデルは真に知的ではない」と主張しています。[24]

知能と意識に関する根本的哲学的見解の対立

この議論の根底には、知能と意識の本質に関する深刻な哲学的対立があります：

計算主義 vs 身体化認知論

計算主義的アプローチ：スケーリング支持派の多くは、知能を情報処理の複雑さと捉え、十分に大規模で複雑な計算システムがあれば意識や理解が創発すると考えています。この観点では、基質（生物学的か人工的か）は重要ではなく、情報処理パターンこそが知能の本質とされます。[25]

身体化認知論：一方、批判派の多くは身体化認知 (Embodied Cognition) の重要性を強調します。この理論によれば、真の知能は物理的身体と環境との相互作用を通じて発達するものであり、純粋な計算では実現できないとされています。Yann LeCunのJEPA理論も、この身体化認知の影響を受けています。[26] [27]

統語論 vs 意味論の問題

統語的处理：現在のLLMは本質的に統語的 (syntactic) な操作、つまりシンボルの操作を行っているに過ぎないという批判があります。John Searleの「中国語の部屋」論証のように、正しい出力を生成できても、真の理解や意味 (semantics) を持っていないとされています。[27]

意味的理解：批判派は、真のAGIには意味的理解、すなわち概念と現実世界の関係を理解する能力が不可欠であると主張します。これには主観的体験や意識が必要であり、純粋な計算的アプローチでは実現できないとされています。[27]

AGIの定義多様性が議論に与える影響

「AGIは達成可能か」という議論の複雑さは、AGI自体の定義が統一されていないことに起因します：

性能ベース定義

OpenAI定義：「経済的に価値のあるほとんどのタスクで人間の能力を上回る高度に自律的なシステム」。この定義は外部から観察可能な性能に焦点を当てています。^[28]

Google DeepMind分類：AGIを5段階（emerging、competent、expert、virtuoso、superhuman）に分類し、competent AGIを「幅広い非物理的タスクで熟練した成人の50%を上回る性能」と定義しています。^[29]

プロセスベース定義

認知アーキテクチャ重視：Song-Chun ZhuのCUVフレームワークでは、AGIを認知アーキテクチャ（C）、理解機能（U）、価値関数（V）の3次元空間の点として定義しています。この定義は内部処理メカニズムを重視します。^[24]

身体化要件：一部の研究者は、真のAGIには物理的環境との相互作用能力（身体化）が不可欠であると主張しています。^[26]

この定義の多様性により、同じ技術的進歩を見ても、研究者によってAGI達成の評価が大きく異なる結果となっています。

現在のAIアプローチを補完・代替する将来的パラダイム

スケーリング仮説の限界を受けて、複数の代替・補完アプローチが提案されています：

世界モデルとJEPA

Yann LeCunのビジョン：物理世界の内部モデルを学習するJEPA（Joint Embedding Predictive Architecture）を提唱しています。このアプローチでは、言語記述ではなく物理観察を通じて世界の動作原理を学習します。^{[16] [17] [18]}

技術的優位性：Meta社のI-JEPAモデルは、16台のA100 GPUを72時間未満使用して6億3200万パラメータの視覚変換器モデルを訓練し、ImageNetで最先端の少数ショット分類性能を達成しています。従来手法と比較して2-10倍少ないGPU時間で同等以上の精度を実現しています。^[17]

ニューロシンボリック AI

ハイブリッドアプローチ：ニューラルネットワークの学習能力と象徴的推論の論理的厳密性を組み合わせることで、現在のAIの限界を克服しようとする試みです。^{[10] [30]}

実装例：物理法則と因果的深層ニューラルネットワークを組み合わせたハイブリッドモデル、情報格子を通じた学習、強化学習技術などが含まれます。

認知的アーキテクチャ

脳インスパイアード設計：人間の認知プロセスを模倣したモジュラー設計により、複数の専門化されたサブシステムの協調を通じて汎用知能を実現しようとするアプローチです。^{[30] [31]}

システム1・システム2統合：Daniel Kahnemanの二重プロセス理論に基づき、直感的（System 1）と分析的（System 2）思考を統合したAIアーキテクチャの開発が進んでいます。^[31]

進化的・神経進化アプローチ

アーキテクチャの進化：勾配降下法に代わり、進化アルゴリズムを使用してニューラルネットワークアーキテクチャを最適化するアプローチです。これにより、従来の手法では発見できない新しいアーキテクチャの創発が期待されています。^[30]

マルチエージェントシステム

専門化と協調：複数の特化したAIエージェントが協調して複雑なタスクを解決するアプローチです。単一の巨大モデルに代わり、タスクに応じて最適化された複数のエージェントの組み合わせにより効率性と適応性の向上を図ります。^[32]

議論の統合的概観と今後の展望

主要論点の整理

合意点：

- 現在のLLMには重要な限界が存在する
- 純粋なスケーリングだけでは不十分である可能性が高い
- 真のAGI実現には新しい技術的ブレークスルーが必要

不一致点：

- スケーリングの有効性継続期間
- 必要な技術的ブレークスルーの性質
- AGI実現の時間軸
- 意識や主観的体験の必要性

現在の技術的収束

興味深いことに、最近の研究動向は両陣営のアプローチが収束する傾向を示しています。OpenAIのo1モデルは推論時計算の拡張を実現し、一方でMeta社のJEPaアプローチは効率的な世界モデル学習を実現しています。これらの進歩は、純粋なスケーリングと代替アプローチの要素を統合したハイブリッド手法の有効性を示唆しています。^{[17] [7]}

投資と研究方向性への示唆

2025年現在、AI業界は転換点にあります。DeepSeekのような効率的なモデルの登場により、単純なスケーリングへの投資効果に疑問が生じています。同時に、より根本的なアプローチへの関心が高まっており、世界モデル、身体化AI、ニューロシンボリック手法への研究投資が増加しています。^[9]

今後5-10年の展望：

- ハイブリッドアプローチの成熟
- 効率性とスケーリングの最適バランス探求

- 評価指標とベンチマークの改善
- 規制フレームワークの発展
- 社会実装における課題の顕在化

結論として、「このままAIが発展するだけでAGI達成可能か」という問いに対する答えは、単純な二分法では表現できない複雑さを持っています。スケーリング仮説は重要な洞察を提供し続ける一方で、その限界も明確になっています。真のブレークスルーは、スケーリングの利点を活用しながら、世界モデル、身体化認知、象徴的推論などの代替アプローチを統合したハイブリッド手法から生まれる可能性が最も高いと考えられます。



1. <https://www.youtube.com/watch?v=2ZLe7KoDLhc>
2. <https://arxiv.org/abs/2001.08361>
3. <https://arxiv.org/pdf/2001.08361.pdf>
4. <https://aisafety.info/questions/7727/Can-we-get-AGI-by-scaling-up-architectures-similar-to-current-ones,-or-are-we-missing-key-insights>
5. <https://dev.to/arbisoftcompany/agi-explained-the-future-of-artificial-intelligence-63f>
6. <https://research.aimultiple.com/artificial-general-intelligence-singularity-timing/>
7. <https://note.com/repkuririn7/n/n33355986f33f>
8. <https://www.newyorker.com/culture/open-questions/what-if-ai-doesnt-get-much-better-than-this>
9. <https://gigazine.net/news/20250328-current-ai-agi/>
10. https://www.reddit.com/r/singularity/comments/1k6pxgu/lms_wont_scale_to_agi_but_instead_well_need/
11. <https://foundationcapital.com/has-ai-scaling-hit-a-limit/>
12. <https://the-decoder.com/sam-altman-on-agi-scaling-large-language-models-is-not-enough/>
13. <https://www.wired.com/story/sam-altman-says-the-gpt-5-haters-got-it-all-wrong/>
14. <https://gwern.net/scaling-hypothesis>
15. <https://www.sprind.org/en/ai>
16. <https://www.onhealthcare.tech/p/beyond-language-models-yann-lecuns>
17. <https://ai.meta.com/blog/yann-lecun-ai-model-i-jepa/>
18. <https://www.linkedin.com/pulse/yann-lecuns-shift-from-lms-jepa-world-models-anshuman-jha-cyyac>
19. <https://jacob buckman.com/2022-06-14-an-actually-good-argument-against-naive-ai-scaling/>
20. <https://garymarcus.substack.com/p/breaking-news-scaling-will-never/comments>
21. <https://garymarcus.substack.com/p/breaking-openais-efforts-at-pure>
22. <https://humancompatible.ai/blog/2023/09/11/stuart-russell-testifies-on-ai-regulation-at-u-s-senate-hearing/>
23. <https://people.eecs.berkeley.edu/~russell/papers/russell-senate23-written.pdf>
24. <https://www.science.org/content/article/ai-gets-mind-its-own>
25. <https://www.linkedin.com/pulse/artificial-intelligences-view-problem-consciousness-michael-watkins-xdkpf>
26. <https://www.zygonjournal.org/article/id/15466/>

27. <https://philarchive.org/archive/MASNQN>
28. <https://openai.com/index/planning-for-agi-and-beyond/>
29. https://en.wikipedia.org/wiki/Artificial_general_intelligence
30. <https://www.nature.com/articles/s41598-025-92190-7>
31. <https://arxiv.org/html/2509.14474v1>
32. <https://www.getdynamiq.ai/post/multi-agent-ai-systems-definition-benefits-limitations-how-to-build>
33. <https://www.aialign.net/blog/transformer-agi-part1>
34. <https://www.pnas.org/doi/10.1073/pnas.2311878121>
35. https://note.com/showying_art/n/n0c5d22aa59e8
36. https://www.reddit.com/r/agi/comments/1msm7mh/current_neural_networks_and_algorithms_are_not/
37. <https://www.scribd.com/document/821176100/Scaling-Laws-for-Neural-Language-Models>
38. <https://xtech.nikkei.com/atcl/nxt/column/18/02423/041800005/>
39. <https://www.docswell.com/s/DeepLearning2023/Z1J7Y2-dlscaling-laws-for-neural-language-models>
40. <https://fortune.com/2025/02/19/generative-ai-scaling-agi-deep-learning/>
41. https://medical-science-labo.jp/scaling_law/
42. <https://www.semanticscholar.org/paper/Scaling-Laws-for-Neural-Language-Models-Kaplan-McCandlish/e6c561d02500b2596a230b341a8eb8b921ca5bf2>
43. <https://www.ibm.com/think/news/agi-right-goal>
44. <https://www.mof.go.jp/pri/research/seminar/fy2025/lm20250625.pdf>
45. https://www.reddit.com/r/singularity/comments/1b96svl/yann_lecun_meta_ai_open_source_limits_of_llms_agi/
46. <https://arxiv.org/html/2502.01677v1>
47. <https://news.ycombinator.com/item?id=40000574>
48. <https://ai.plainenglish.io/experts-weigh-in-on-whether-the-current-transformer-architecture-will-get-us-to-agi-c32aa55c3ed6>
49. <https://www.ml-science.com/blog/2024/10/10/the-path-to-artificial-general-intelligence-yann-lecuns-vision-for-the-future>
50. <https://openreview.net/forum?id=vMTijVnXQ8>
51. <https://lexfridman.com/yann-lecun-3-transcript/>
52. https://www.reddit.com/r/BetterOffline/comments/1mg1y4m/what_do_you_all_think_of_gary_marcus_hes_been/
53. https://www.linkedin.com/posts/yann-lecun_a-good-article-by-cade-metz-at-the-nyt-about-activity-7329654030316548096-C8Nm
54. <https://www.linkedin.com/pulse/beyond-transformers-next-ai-architectures-siddharth-bhalsod-bdmkc>
55. <https://toloka.ai/blog/agi-vs-other-ai/>
56. <https://www.ibm.com/think/topics/artificial-general-intelligence>
57. <https://time.com/6556168/when-ai-outsmart-humans/>
58. <https://www.quantamagazine.org/artificial-intelligence-aligned-with-human-values-qa-with-stuart-russell-20150421/>
59. <https://www.coursera.org/articles/what-is-agi-versus-ai>

60. <https://aws.amazon.com/what-is/artificial-general-intelligence/>
61. <https://www.linkedin.com/pulse/agi-revolution-how-close-we-achieving-human-level-ai-hugo-raaijmakers-ef5ne>
62. <https://kanerika.com/blogs/ai-vs-agi-vs-asi/>
63. <https://arxiv.org/pdf/2505.11866.pdf>
64. <https://www.mckinsey.com/featured-insights/mckinsey-explainers/what-is-artificial-general-intelligence-agi>
65. https://www.reddit.com/r/singularity/comments/17xzg0d/the_agi_hypothesis_for_why_sam_altman_was_ousted/
66. <https://bito.ai/blog/ai-agent-vs-llm/>
67. <https://docs.uptiq.ai/overview-of-genai/key-concepts/ai-agents-vs-llm-based-apps>
68. <https://deepfa.ir/en/blog/world-model-ai-agi-future>
69. <https://www.anubavam.com/blog/ai-vs-ai-agents-llms-systems>
70. https://www.reddit.com/r/philosophy/comments/vcrds5/the_hard_problem_of_ai_consciousness_the_problem/
71. <https://rohitbandaru.github.io/blog/JEPA-Deep-Dive/>
72. <https://arxiv.org/html/2508.17281v1>
73. <https://www.robometricsagi.com/blog/ai-policy/embodiment-experiences-feeling-consciousness>
74. <https://www.youtube.com/watch?v=yUmDRxV0krg>
75. <https://www.ibm.com/think/topics/ai-agents-vs-ai-assistants>
76. <https://deepsquare.jp/2022/08/autonomous-machine-intelligence/>