

「AGIはもう来ている」 の解剖学

量子物理学者の発言から読み解く、AI研究の
「真のパラダイムシフト」と「ハイプ」の境界線

マクシメンコ氏の実際のインタビュー発言

“Guys, I think AGI is here.”

"It's like a weak AGI in some sense."

- ⊕ Requires correct prompting, expertise, and guidance.
- ⊕ Not fully self-consistent.

Forbes JAPAN / Yahoo! ニュース見出し

WARNING:
CONTEXT REMOVED

「汎用人工知能・AGIはもう来ている」

本人の自己評価は「専門家の誘導が必要な弱いAGI (Weak AGI) 」。
完全な自律性を持つ一般的なAGIの宣言ではない。

truth

検証ダッシュボード：シグナルとノイズ



LOW

「AGIがすでに到来した」という証明

一般性、自律性、頑健性が未検証。科学的証拠として不十分。



MEDIUM

現在のNISQ量子計算 + AIの有用性

有望だがタスク依存。古典アルゴリズムの進歩により「量子優位」の境界は常に変動する。



HIGH

AIエージェントによる科学研究の劇的な高速化

マクシメンコ氏以外 (Google, DeepMind等) の事例でも実証済み。これが真のパラダイムシフト。

AGIの定義マトリクス：何を以て「到来」とするか

	汎用性	自律性	対象タスク	本件での立証度
OpenAI Charter (厳格)	最も経済的価値のある仕事	高度な自律性	人間を上回る	✗ 未立証
DeepMind Level 1-2 (段階的)	幅広い非物理 タスク	非熟練者～ 成人50%水準	段階的評価	△ 証明不足
ARC-AGI (スキル獲得)	未知の環境での 一般化	人間に匹敵する 効率	探索と計画実行	✗ 未立証
マクシメンコの "Weak AGI" (私的用法)	限定的な高度研究	専門家によるプロンプトと誘導が必須	広範な研究実行	✓ 成立し得る

記事の「AGI」は学術的な万能知能ではなく、限定的な「高度なツール化」を指す私的分類である。

検証1：「半年の研究を1週間で」の抜け穴

人間の博士課程 - 6ヶ月

問題設定の発明

失敗と探索

ソフトウェア実装

結果解釈

AIによる再構築 - 1週間

実装

数式のバグ発見

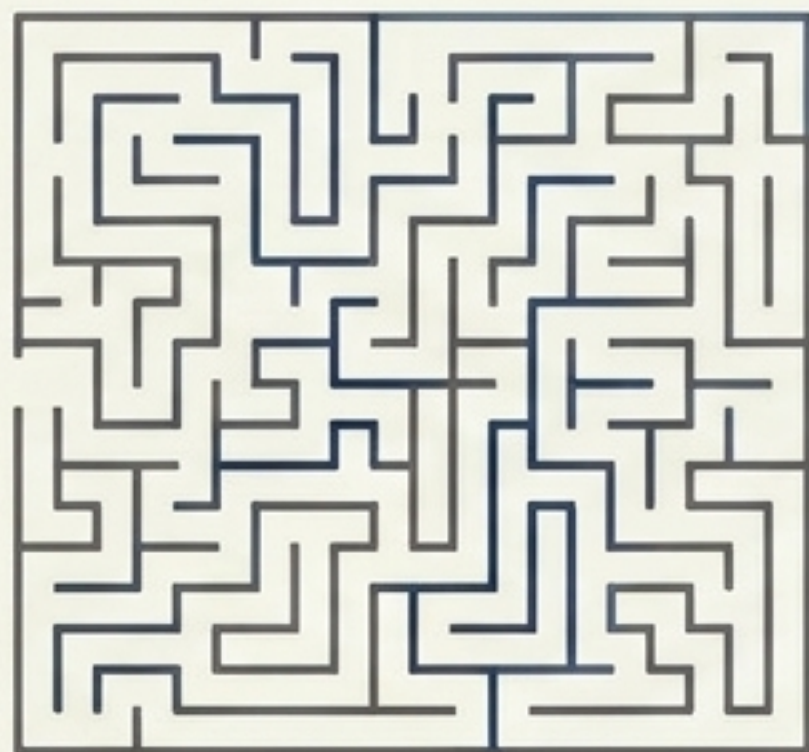
データ汚染の疑い

本件は「既知研究の再構築」か「未知の発見」か？

過去論文がLLMの訓練データに含まれていたか (Retrieval Leakage) を排除するブラインドテストがない限り、科学的証拠とはならない。

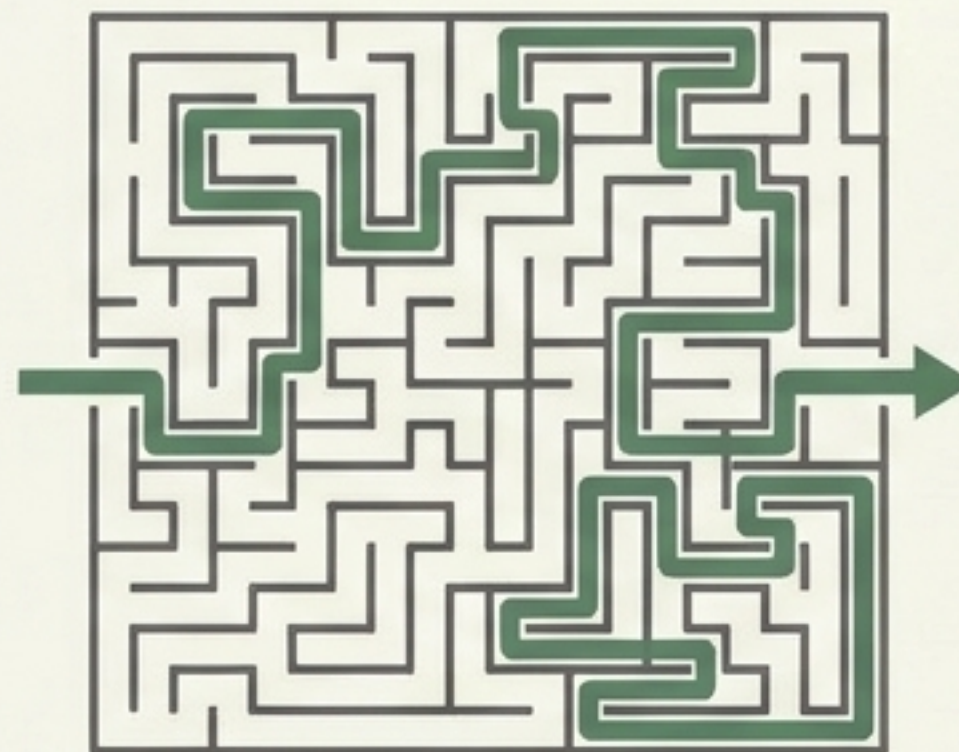
検証2：情報状態の非対称性

オリジナルのQ4Bioチーム（数年の研究）



暗闇の探索空間。何を試すべきか、どの問題設定が正しいかを発明する過程。

AIエージェントによる再現（数日）

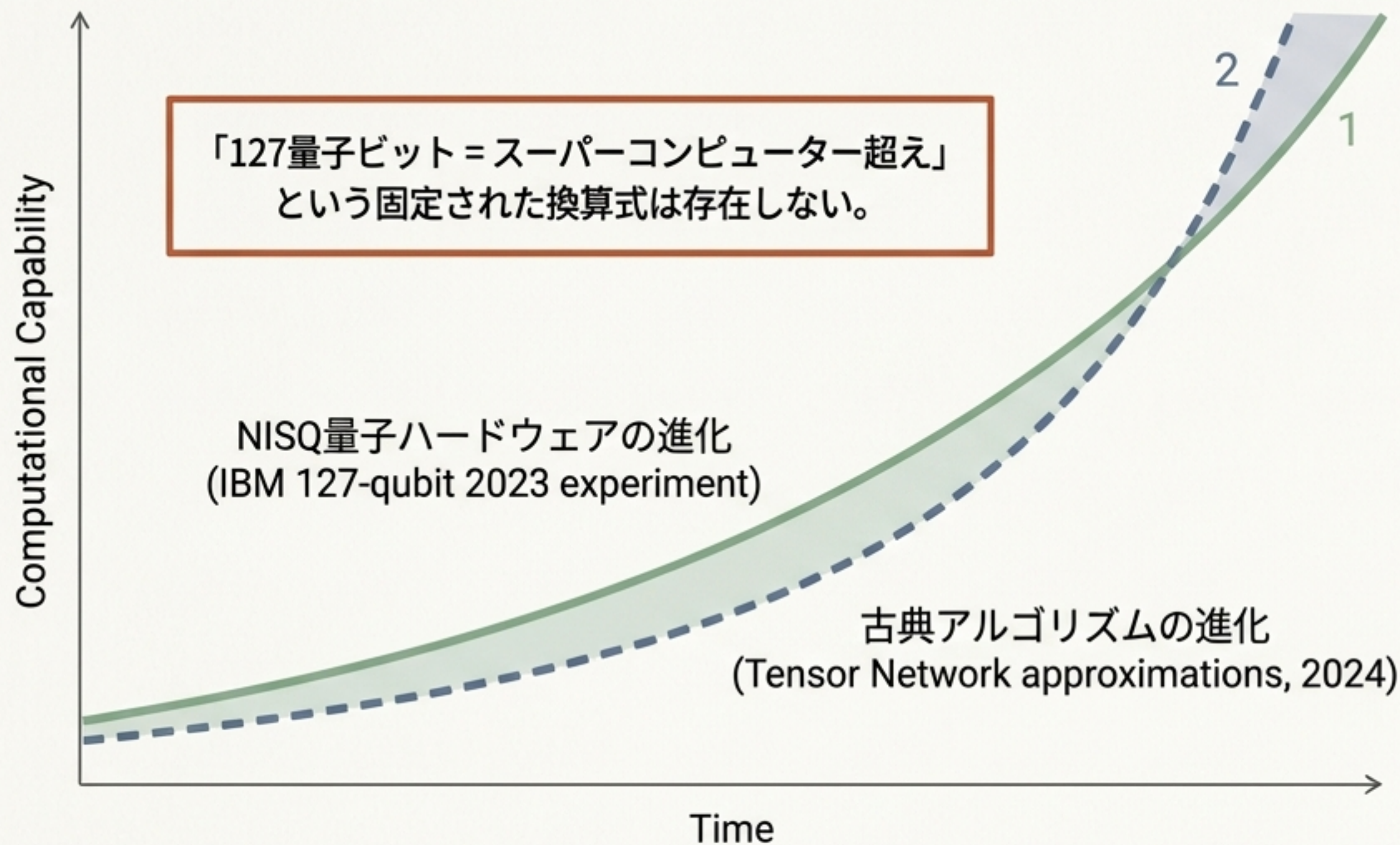


地図がある道。すでに「成功例が公表された後」の再実装。

後知恵バイアス

「完全なゲノムを量子状態としてロードできたこと（事実）」を、「AIが未知の課題を自律発見できること（AGIの証明）」と混同してはならない。

検証3：「量子Utility」の動的境界線



逃げ水現象

古典計算側（テンソルネットワークなど）も進歩するため、量子優位の境界線は常に逃げていく。量子ビット数だけでAGI到来の証拠とするのは不正確。

検証4：ベンダー内部ベンチマークの危うさ

“9時間・30,000ドルの分子動力学シミュレーションが、約30秒・25ドルに？”

独立検証に必要な要件:

- ☒ Haiqu側からの製品・内部テスト発表
- ☐ 元の30,000ドル実装の最適化レベルの公開
- ☐ 同一の数値精度の担保
- ☐ クラウド料金・ハードウェア条件の完全一致
- ☐ 査読論文による第三者の独立再現

結論：現段階では「優れたベンダーの宣伝値」であり、独立した科学的事実として扱うべきではない。

真のシグナル：AIによる科学研究の「圧縮」は実在する

Google Research “AI co-scientist”

LLMエージェントが公開文献のみを用いて、未公開の抗菌薬耐性研究を独立再発見。

DeepMind “AlphaEvolve”

50以上の数学問題で既知の最良解を再発見し、量子物理用回路のエラーを従来の1/10に抑える設計を自律探索。

共通項

これらは「完全自律型AGI」ではなく、人間が目標を設定する「Scientist-in-the-loop（専門家介在型）」の複合システムである。

再構築：新たなパラダイムの方程式

~~[単独のAI] ≡ [汎用的な人間の研究者 (AGI)]~~

[人間の専門家]

+

[フロントティアLLM]

+

[エージェント・
スキャフォールディング]

+

[量子計算等の
専門ツール]

=

弱いAGIのように見えるほど強力なオーケストレーション

量子コンピューターは「AGIを生み出す魔法の箱」ではなく、この複合システムの中でLLMが操作する「極めて強力な計算ツール」として機能する。

The Verification Blueprint: 次の「ハイプ」を評価する3つの判断軸

1. 成功率と失敗の開示

「1回の奇跡的な成功」か、「安定したタスク遂行」か？
事後選択バイアスを疑え。

2. 情報の非対称性の排除

公開済みデータの「再現」か、未知領域での「新発見」か？
ブラインドテストの有無を確認せよ。

3. ツールと知能の分離

「AIそのものが賢くなった」のか、「AIが既存の高度なツール（量子など）をうまく使えるようになった」のか？

「AGI到来」というノイズに惑わされず、
「複合研究システムによる時間の圧縮」という真のシグナルを見極めること。