

フランソワ・ショレ「AGIへの道」講演の要約と詳細レポート

要約

- ・現在のディープラーニングのスケール路線は限界に達しており、これ以上モデルを巨大化しても**真の汎用人工知能（AGI）には到達できない**とショレは主張しています¹。現行のAIは訓練データのパターン記憶に頼って未知の問題へ柔軟に対応できないためです²。
- ・ショレは「**知能**」とは**未知の問題を解決できる流動的な適応能力**であると捉えており、従来のベンチマークはこの能力を正しく測れていないと指摘します³。例えば彼の提唱した抽象的推論コーパス（ARC）では、人間が95%以上正解できる課題で最先端モデル（GPT-4.5相当）ですら正解率は約10%に留まり、巨大モデルの大規模化だけでは柔軟な知性が得られないことを示しています⁴。
- ・AGI実現への道筋として、ショレは**深層学習による直観的パターン認識と、プログラム合成による論理的推論を組み合わせたハイブリッドアプローチ**を提唱しています⁵。ディープラーニングで獲得した知識（機能モジュール）を基に、その場で新しいプログラムを合成・実行することで、未知のタスクにも柔軟に対処できるAIを目指すものです⁶。
- ・最終的なビジョンは、**人間のプログラマのように問題に応じて自律的にアルゴリズム（プログラム）を組み立てるメタ学習型のAIを実現すること**です⁷。ショレはモデルが推論時に自己を書き換えて新規タスクに適応する「テスト時適応（TTA）」の重要性を強調しており⁸、このアプローチによってAIを熟練した人間のプログラマと同等に柔軟かつ創造的な存在に高めたいと考えています⁹。

ショレのAGIに関する核心的主張とビジョン

ショレは現在主流となっている大規模言語モデル（LLM）を中心としたディープラーニングの延長では**真の汎用人工知能には到達できない**と明言しています¹。モデルの規模拡大によって一時的な性能向上は見られるものの、それは膨大なデータからパターンを記憶しているに過ぎず、新たな問題に直面した際に対応できる「**適応力**」を欠いているためです¹⁰。実際、AI研究者の多数（約76%）も「現在の手法をスケールアップするだけではAGI実現は困難」と考えており、この点でショレの見解は専門家の間でも広く共有されています¹¹。

ショレは「**知能**」とは**過去に経験のない問題を解決できる能力である**というビジョンを掲げています。彼はマービン・ミンスキーの「人間が既にできる作業を自動化するのが知能」という見方ではなく、ジョン・マッカーシーの「**未知の問題を解く能力こそ知能**」という定義を支持し³、「**スキル**」と「**知能**」を明確に区別しています。彼自身の比喻では、スキルとは既存の道路を速く走れること、知能とは**まだ誰も走ったことのない未知の地に新しい道を切り拓く力**だと説明しています¹²。このように**未知への適応と新規解決策の創出こそがAGIの本質**であり、ショレのビジョンにおける汎用知能の指標となっています。

ショレが提唱するAGIへのアプローチと技術的視点

そうしたビジョンに基づき、ショレは**ディープラーニングとプログラム合成を組み合わせたハイブリッドアプローチ**を提唱しています⁵。現在の深層学習モデルは大量データからパターンを認識する直観的能力に優れていますが、一方で厳密なルールに基づく推論やシンボル操作など**論理的思考を要する課題では性能が極端に低下する**ことが指摘されています¹³。ショレはこの弱点を補うために、ディープラーニングによって得られた**知識の断片（機能モジュール）**を活用し、それらを組み合わせて新しいプログラム（解法）をその場で合成する仕組みを導入すべきだと主張します⁶。これにより、システムは**学習時に記憶したスキルを部品と**

して再利用しながら、直面した課題ごとに新たなアルゴリズムを構築できるようになります。すなわち深層学習で培った直観とプログラム合成による論理的探索を融合させることで、既存の知識の単なる適用を超えて真に新規な問題解決が可能になるというアプローチです¹⁴。

この技術的視点の要となるのが「テスト時適応 (test-time adaptation)」です。ショレは2024年頃から研究分野で注目が高まったこの潮流に言及し、**モデルが推論（テスト）時に自らの内部状態や振る舞いを変化させて問題に取り組む方法論**がAGIへのブレークスルーにつながると述べています⁸。具体的な手法として彼が挙げるのが**プログラム合成**や**チェイン・オブ・ソウト（思考過程の列）**の生成です¹⁵。モデルは与えられた課題に対し、単一の一発推論で答えを出すのではなく、自ら一連のステップ（思考過程）を試行錯誤して組み立て、その手順に従って解を導きます。実際、このような**深層学習によるガイド付きプログラム探索**が威力を発揮し始めており、ショレによれば特殊な手法を組み込んだOpenAIの「o3」モデルはARCベンチマークで人間に匹敵する成績を初めて達成しました¹⁶。この成果は**適応的なプログラム生成というアプローチが、汎用的な問題解決能力を飛躍的に高めることを示す具体例**と言え、ショレの主張する方向性を裏付けています。

現在のディープラーニング/AI研究の限界に対する見解

図: ARC（抽象的推論コーパス）のタスク例。 人間には直感的に解ける視覚パズルですが、学習済みAIにとっては極めて困難です。左端から複数組の入力グリッドとそれに対応する出力グリッドが提示されており、最右列のグリッドはAIが解答すべき新たな入力（出力は空欄）を示しています。ARCはこのような「**人間には容易だがAIには困難**」な課題を通じてAIの適応力の限界を浮き彫りにする指標となっています¹⁷。実際、ARCでは人間がほぼ100%正解できる問題でも、最新の大規模モデルで正解率はせいぜい一桁台（GPT-4.5で約10%）に留まります⁴。これは現在の深層学習システムが直面する**一般化能力の不足**を如実に物語っています。

ショレは、現在のAIが抱える限界として「**データから暗記したスキルの範囲内でしか対処できない**」点を強調します。大量の事前学習により既知のパターンには高い性能を発揮するものの、一旦トレーニング分布から外れた事例や初めて見るタイプの問題に対しては極端に脆弱であるという弱点です¹⁸。彼が設計したARCベンチマークはまさにこの点を突くものであり、既存のAIが**いかに訓練データに依存した「結晶化した知識」**に留まっているかを浮き彫りにしました¹⁹。ARCは人間の幼児でも解けるような常識的パターンを含む問題にもかかわらず、汎用性のないAIは対応できず、高度に特化した現在のAI技術の限界を示しています。

さらにショレは、近年LLMの能力向上が報告されている領域の多くが**人手による注釈付けやタスク固有の微調整に支えられた「付け焼き刃の改善」**に過ぎない点を批判しています²⁰。例えばモデルが論理的推論を苦手とする問題に対し、開発者側がチェイン・オブ・ソウト（思考過程）の手順を人為的に教え込んだり、追加の訓練データを与えたりすることで一見性能が上がる場合があります。しかしショレによれば、こうした**人間による場当たりの弱点補強は本質的解決にならず、新たな状況に対する汎用的適応力が向上したわけではないのです**²⁰。このような手法は労力も大きく、無数に存在し得る未知のケース全てに対処することは不可能です。したがって、現在のディープラーニング中心のアプローチには構造的な限界があり、**真の汎用性を得るには根本的なアーキテクチャの再考が必要だ**とショレは述べています。

論理的・哲学的・倫理的含意

ショレの提起した議論は、AIにおける知性の捉え方を巡る哲学的含意も含んでいます。彼が強調する「**未知への適応こそ知能**」という見解は、既知のタスクを自動化することに注力してきた従来のAI研究の姿勢を見直すものです³。ミンスキー的な「タスクの自動化から知能を定義する」路線ではなく、マッカーシー的な「**新たな問題への対応力**」で知能を測ろうというアプローチは、**AI研究の評価基準や目標設定の哲学的転換**を促すものと言えます¹²。これにより、今後のAIベンチマークは単に既定の課題で高得点を競うのではなく、未知の課題での創造的問題解決能力をいかに引き出すかに焦点を当てるべきだという示唆が得られます。

倫理的観点について、ショレは本講演で直接的な言及はしていません。彼の関心はもっぱら技術的ロードマップと知能の定義にあり、AIの安全性や社会への影響といった課題は深く掘り下げられませんでした。しかし、そのビジョンには暗黙の倫理的含意も読み取れます。ショレが目指す「**自ら発明し適応できるAI**」は、正しく運用すれば人類の科学的発見や問題解決を加速させるポジティブな力になると考えられます²¹。一方で、AIが自律的にプログラムを生成・実行するという方向性は、制御性や透明性の確保といった課題も伴うでしょう。ショレ自身は講演内でこれらリスク面には触れていませんが、提案されたアプローチを進める上では**AIが意図せぬ暴走をしないよう管理する仕組みや人間との目標の整合性（アラインメント）**についても今後検討が必要となることは言うまでもありません。総じて、ショレの議論は技術的展望に集中しているものの、それは同時に「**知能とは何か**」「**汎用AIをどう実現し何に活かすか**」という根源的・倫理的問いを投げかける内容でもあります。

発言の文脈と背景

フランソワ・ショレはフランス出身のソフトウェアエンジニア・研究者で、深層学習ライブラリ「Keras（ケラス）」の作者として知られます。彼は約9年間にわたりGoogleでディープラーニングの研究開発に従事した後、**真の汎用AI実現に向けた独自のアプローチを追求するため自らAI研究ラボを立ち上げました**²²。2019年には論文「On the Measure of Intelligence（知能の測定について）」を発表し、**人間並みの柔軟な問題解決力を測る新たなベンチマーク**として抽象的推論コーパス（ARC）を提唱しています²³²⁴。この背景には、当時すでにディープラーニングが画像認識やゲーム攻略など限定的なタスクで顕著な成果を出す一方、領域を超えた汎用性や初見の課題への対応力に課題があると認識され始めていたことがあります。ショレのARCはまさにそのギャップを定量化する試みであり、初回のARC競技会（2020年）で優勝チームの正解率が21%程度に留まった結果²⁵²³は、彼の指摘する現在のAIの限界を裏付ける証拠となりました。

本講演「AGIへの道（How We Get To AGI）」は、2024年8月にシアトルで開催された国際会議 AGI-24 における基調講演として行われたものです²⁶。著名な研究者によるこの講演は大きな注目を集め、ショレの見解が現時点でのAI研究コミュニティにおいて重要な論点であることを示しています。講演後、彼は自身のビジョンを実現すべくさらなる行動を起こしています。ARCを発展させた「**ARC Prize**」というコンペティションやベンチマークの整備に加え、2025年には新たな研究ラボ「NDEA（エヌデア）」を創設し、提唱するアプローチ（深層学習＋プログラム合成）による次世代AIの開発に乗り出しました²¹。NDEAでは「**発明・適応・革新**」をキーワードに、**人間と同程度に効率よく学習し続けるAI**の研究開発を目指しているとされます。こうした経歴と取り組みからも分かる通り、ショレの発言は単なる理論的提案に留まらず、**自身の豊富な経験と実績、そして現在進行形のプロジェクトに裏打ちされた実践的ビジョン**だと言えるでしょう。²²²³

1 2 3 4 7 8 9 10 12 13 14 15 16 21 François Chollet on the end of scaling, ARC-3 and his path to AGI

<https://the-decoder.com/francois-chollet-on-the-end-of-scaling-arc-3-and-his-path-to-agi/>

5 6 11 18 20 22 26 François Chollet on why LLMs won't scale to AGI — EA Forum

<https://forum.effectivealtruism.org/posts/MGpJpN3mELxwyfv8t/francois-chollet-on-why-llms-won-t-scale-to-agi>

17 19 23 24 25 ARC Prize - What is ARC-AGI?

<https://arcprize.org/arc-agi>