

News&Trend

IBM、  
積層チップでムーアの法則を10年延命  
「風邪はしょうがない」を覆せるか

Interview

野入亮人 (理化学研究所)

# MIT Technology Review

Published by KADOKAWA / ASCII

Vol.

87

2026.07



AIの10大潮流 [2026年版]

003

特集

2026.07

## AIの10大潮流 [2026年版]

- 01 ヒューマノイド・データ
- 02 LLM+
- 03 AI強化型詐欺
- 04 世界モデル
- 05 新・作戦司令室
- 06 武器化されたディープフェイク
- 07 エージェント・オーケストレーション
- 08 オープンLLM
- 09 人工科学者
- 10 AI抵抗運動

024 AI閉塞感の時代、  
私たちはまだ何も分かっていない

027 崖っぷちのプライバシー、  
AIエージェントが崩す最後の砦

---

031 U35 イノベーターの軌跡 #39  
野入亮人（理化学研究所）  
シリコン量子ビットで大規模化の壁に挑む研究者

034 News&Trends  
IBMが積層チップ、「ムーアの法則」10年延命へ  
アンソロピック、創薬に照準 「Claude Science」を発表  
「固体冷却」に熱視線、冷房市場に割って入れるか  
「風邪はしょうがない」覆せるか 民間主導の呼吸器ウイルス根絶計画  
成層圏に浮かぶ5G基地局が この夏、日本にやってくる

# AIの 10大潮流

[2026年]

人工知能 (AI) をめぐる喧騒の中で、本当に目を向けるべきものは何か。AIの進化を追い、その次に来るものを予測してきたMITテクノロジーレビューは、2026年上半期時点の10大潮流をリスト化して発表する。進化の壁に直面する大規模言語モデル (LLM)、人型ロボットや最先端のエージェント技術、科学の自動化は私たちに何をもたらすのか。AI覇権を争う米中の動向や、兵器化するディープフェイク、巧妙化する詐欺にどう立ち向かうべきなのか。浮かび上がったのは、AIが紡ぐすばらしい未来と恐ろしい現実の姿だ。

## INDEX

- 01 ヒューマノイド・データ
- 02 LLM+
- 03 AI強化型詐欺
- 04 世界モデル
- 05 新・作戦司令室
- 06 武器化されたディープフェイク
- 07 エージェント・オーケストレーション
- 08 オープンLLM
- 09 人工科学者
- 10 AI抵抗運動

Congra  
you ha  
To che  
please

# 01 ヒューマノイド・ データ

## AIロボブームで生まれた 奇妙な労働

ロボティクス企業は今や、ヒューマノイドの訓練データとして、ギグワーカーたちにカメラやセンサーを装着させて人間の手や手足の動かし方に関する膨大なデータを収集している。ただし、この奇妙な労働がブレイクスルーをもたらすかどうかは不明だ。

**最** 近、食べ物をボウルに入れ、電子レンジで温め、取り出すといった作業を動画に撮ると暗号通貨が支払われるというアプリへの参加を勧められた。また別のWebサイトでは、中国・深センにあるロボットアームを遠隔操作してパズルや課題をこなす新しいゲームを試してみないかと提案された。ロボットの器用さを向上させるためだという。

一体何が起きているのか。ちょうど私たちの言葉が大規模言語モデル（LLM）の訓練データになったように、ロボティクス企業は今、人間の動作に関するデータが高性能なヒューマノイド（人型ロボット）の開発に役立つと見込んでいる。ヒューマノイドは単純なロボットアームよりも訓練が難しいにもかかわらず、現在人間が働く場所により容易に溶け込める（そしていつかは人間を完全に置き換える）と各社は考えている。

ヒューマノイドの訓練方法に関するこの新たな

発想は、2022年のChatGPT（チャットGPT）の登場に端を発すると言える。LLMは、人工知能（AI）企業が収集できた（あるいは盗んだと主張する声もある）あらゆる文章という膨大な訓練データにさらされることでテキストを生成できるようになった。ロボット研究者たちはこのスケーリング則をロボット工学に応用したいと考えたが、人間の動作を記述したインターネット規模のデータコレクションが存在しなかった。

そのようなデータを大量に収集することの困難さに直面した企業は、バーチャル・シミュレーション内でロボットに動作を学習させるといった代替手段を用いた。しかし、シミュレーションは摩擦や弾性といった現実世界の物理現象を完全には再現できないため、シミュレーション内で訓練したロボットは文字通りつまづく傾向があった。

今やヒューマノイドを開発する企業は、手間がかかるとはいえ、現実世界のデータを収集するこ

## ヒューマノイド・データ

とで大きな成果が得られると判断するようになった。そこから事態は奇妙な様相を呈し始めた。

初期の取り組みは素朴で学術的なものだった。研究室では、カメラやハンドヘルドグリッパーを装着した人々がワッフルを焼いたり机を片付けたりといった家事をこなす様子を何時間も収録し、そのデータをオープンに共有していた。しかしロボット工学へのベンチャーキャピタル投資が急増し、2025年にはヒューマノイド分野だけで61億ドルにも達した結果、訓練データの作成をめぐる

競争はより熾烈かつ精巧なものになってきた。

現在、中国には訓練センターが存在し、そこでは人々がエクソスケルトン（外骨格、ここでは着用型ロボットを指す）とVR（実質現実）機器を装着してテーブルを拭くといった同じ反復作業を1日に何百回もこなしている。ナイジェリア、アルゼンチン、インドのギグワーカーたちは自宅での家事を動画に撮影している。2026年に入って筆者は、米国のある配送会社が従業員にセンサーを装着させ、荷物を運ぶ際の動きを追跡していることを知った。その目的の1つは労働災害の研究だが、もう1つの目的は従業員を置き換えるロボットの訓練データの収集だという。

こうした動向は、肉体労働者がデータ収集者へと変わっていくという超現実的な労働の未来を示唆している。しかし、収集した動作データでロボットを訓練することは依然として複雑な問題をはらんでいる。技術的なブレークスルーをもたらすために必要とされる規模でそれを実現できるかどうかさえ不明であり、収益性のあるビジネスを構築できるかどうかはなおさら不透明だ。

電子レンジを開ける私の動画には、どれほどの価値があるのか。ロボットに夕食の調理を教えるには、そのような瞬間が何千回分必要なのか。おそらく2026年こそ、その答えが明らかになるだろう。

by James O'Donnell  
(米国版 AI／ハードウェア担当記者)



## 02

## LLM+

## 大量かつ複雑なタスクを、 安く効率よくこなす

大規模言語モデル（LLM）は進化し続けており、いずれは人間が解くのに数日から数週間かかるような複雑かつ多段階の問題を処理できるようになるかもしれない。それを実現するためには、LLMは今よりも低コストで高効率になる必要がある。

2 022年末に実験的なプロトタイプとして登場したオープンAI（OpenAI）のChatGPT（チャットGPT）は、数億人にとって日常的な万能アプリとなった。ChatGPTのような大規模言語モデル（LLM）は新たな未来の象徴となり、テクノロジー業界全体がその熱狂に飲み込まれ、各社は競合製品の開発に躍起になった。

旧世代の熱狂の余韻が冷めやらないにもかかわらず、「次に来るものは何か」という問いを止める者はいない。先に答えを言ってしまうと、LLMの次に来る大きなものは、さらに進化したLLMだ。それを「LLM+」と呼ぶことにしよう。

課題は、人間が解くのに数日から数週間かかるような複雑かつ多段階の問題をLLMに処理させることだ。トップ研究機関が掲げる目標どおり、LLMが人類の最も困難な課題の解決に貢献するためには、より長い時間にわたって自律的に動作できる能力が必要となる。

その実現に向けて、いくつかの進展が求められる。まず、LLMはより効率的かつ低コストで動作できるようにならなければならない。最大の進

歩のいくつかはこの分野で起きている。「専門家の混合（MoE：Mixture of Experts）」と呼ばれるアプローチは、LLMを複数の小さな部分に分割し、それぞれに異なるタスクの専門性を持たせるものだ。これにより、特定の時点ではモデルの一部だけを起動すれば済むようになる。

LLMをより効率化するもう1つの方法として、現在ほぼすべてのLLMの基盤となっているニューラル・ネットワークの一種であるトランスフォーマーを廃止し、画像や動画生成により一般的に使われる別種のニューラル・ネットワークである拡散モデルに切り替えるという案もある。さらに実験的なアプローチも存在する。2025年に中国の人工知能（AI）企業ディープシーク（DeepSeek）はテキストを画像にエンコードする手法を発表し、計算コストの削減を実現した。

もう1つの重要な進歩の領域は、LLMの「コンテキスト・ウィンドウ」と呼ばれるものに関係している。これはモデルが一度に取り込めるテキスト（または動画）の量であり、作業記憶に相当する。数年前、LLMが一度に処理できるのは数千トークン（単語または単語の一部）、すなわち数十ペ

ージ分のテキストにすぎなかった。最新モデルのコンテキスト・ウィンドウは最大100万トークンに達し、書籍を何冊も積み重ねた分量に相当する。しかし、コンテキスト・ウィンドウが大きくなり、タスクが長くなるほど、モデルが脱線したり、作業内容を忘れていたりする可能性が高まる。

この分野でもブレイクスルーが起きている。MIT（マサチューセッツ工科大学）コンピュータ科学・人工知能研究所（MIT CSAIL）の研究者らが最近発表した論文では、「再帰的LLM（recursive LLMs）」と呼ばれる手法が提案されている。広大なコンテキスト・ウィンドウを一度に取り込む代わりに、再帰的LLMは入力をチャンク（塊）に分割し、各チャンクを自身のコピーに送る。そのコピーはさらにチャンクを細分化し、結果をさらに多くのコピーに送ることもある。より小さな情報の断片を複数のLLMが並列処理することで、長く困難なタスクに対してはるかに高い信頼性が得られるようだ。その結果は確かにLLMではあるが、これまで知られていたものとは異なる。

by Will Douglas Heaven  
(米国版 AI担当上級編集者)



Stephanie Arnett/MIT Technology Review | Adobe Stock

Insider Online限定

eムックはMITテクノロジーレビュー[日本版]の  
有料会員限定サービスです。  
有料会員はすべてのページ、バックナンバーを  
ダウンロードできます。

ご購入はこちら



<https://www.technologyreview.jp/insider/pricing/>

No part of this issue may be produced by any mechanical, photographic or electronic process, or in the form of a phonographic recording, nor may it be stored in a retrieval system, transmitted or otherwise copied for public or private use without written permission of KADOKAWA ASCII Research Laboratories, Inc.

本書のいかなる部分も、法令または利用規約に定めのある場合あるいは株式会社角川アスキー総合研究所の書面による許可がある場合を除いて、電子的、光学的、機械的処理によって、あるいは口述記録の形態によっても、製品にしたり、公衆向けか個人用かに関わらず送信したり複製したりすることはできません。