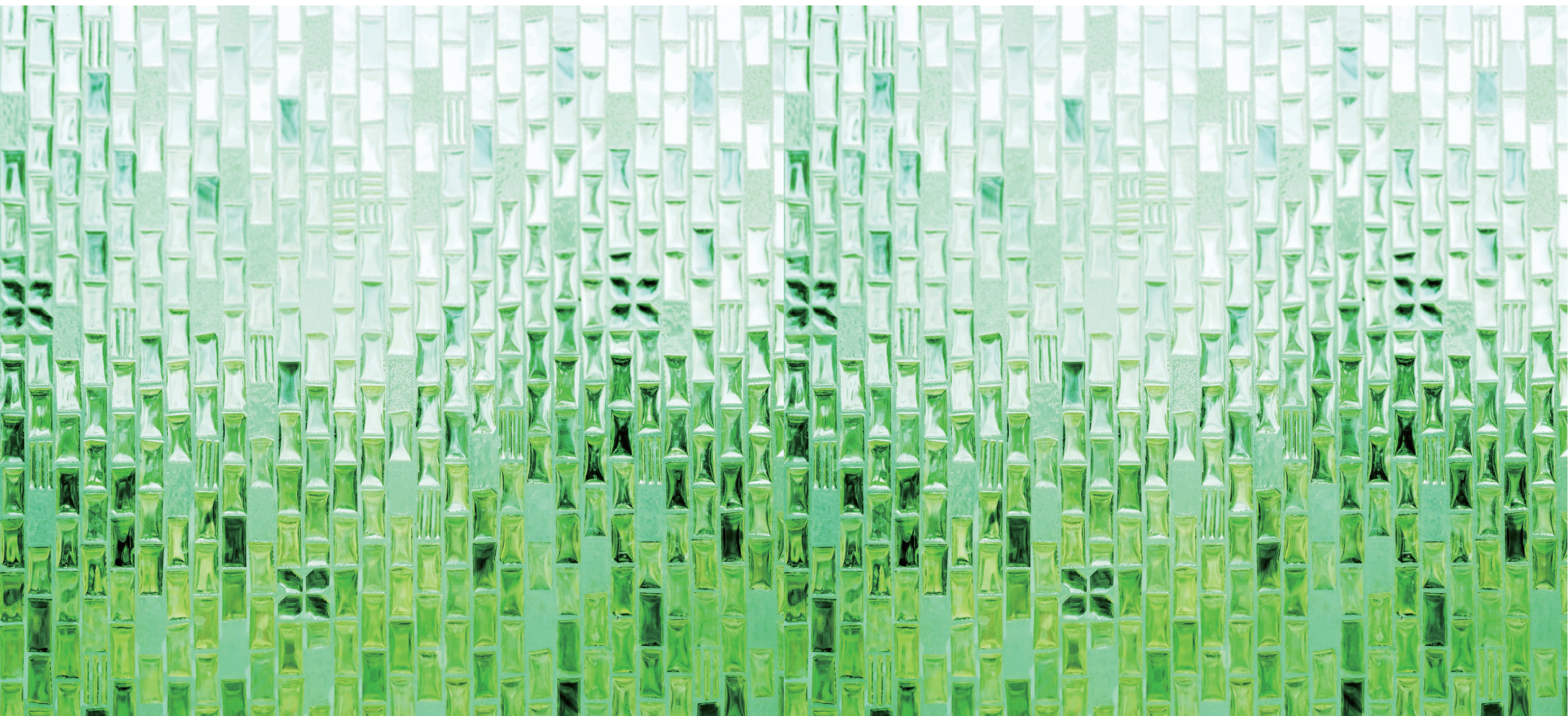


特集 AIの歴史と未来



特集 AIの歴史と未来

急速な発展・普及を遂げ、
社会全体に大きな影響を与えているAI。
すでに現代における必須の技術となっているが、
AIとは何で、どのような可能性を秘めているのか。
AIが形づくる未来のビジョンに肉薄する。

〃言語理解〃というAIの本質

——近年、日常生活からビジネス領域まで、あらゆるシーンに影響を与えているAIですが、そもそも、AIとはどういう存在だと定義つけられますか。

黒橋 何をもってAIと呼ぶか、その定義は簡単ではありません。というのも、ChatGPTやGeminiのように、人間の知的活動をサポートするような「本当の賢さ」をAIが持ち始めたのはごく最近のことで、それまでは、AIに対する認識はかなり曖昧なものだったからです。

AIの賢さというのは、例えばスーパーコンピュータの計算能力のような、量的な賢さとは異なります。もちろん、計算速度やそのベ-

AIを受容した先に見える未来

国立情報学研究所 所長
大規模言語モデル研究開発センター センター長

黒橋 禎夫

生活の中ですでに身近な存在となったAI。情報技術が目覚ましい発展を遂げていくなか、私たちは今後、どのようにAIと向き合い、そのテクノロジーを活用していくべきだろうか。大規模言語モデルを研究する黒橋禎夫氏に、AIがもたらす社会への影響と展望を聞いた。

スとなる半導体技術など、情報開発全体の進歩がAIを支えているのは間違いありません。ただ、AIに求められるのは、様々なメディア情報や人間の言語など、複雑な文脈を柔軟に解釈し表現するといった、言わば質的な賢さ。そうした意味では、AIの一番の本質は「言語を理解すること」だと、私は考えています。それはすなわち、膨大なテキストデータをもとに人間の言語をより自然に理解・生成する能力を持ち、人間とAIとの対話を可能にした「大規模言語モデル（LLM）」のことです。

——そうした質的な賢さを備えたAIが、これまでどのように進化してきたのか教えてください。

黒橋 〃AI (Artificial Intelligence)〃という言葉は、1956年にアメリカの大学で開催

された、人工知能に関する研究発表の場であるダートマス会議において、初めて使用されました。この会議で、人間の思考を機械で再現する「人工知能（AI）」という概念が確立されたのです。これ以降、AIという言葉が認知され、研究開発は活発化しました。これが、いわゆる第一次AIブームです。しかし、当時のコンピュータの性能には限界があり、人間の知能のモデル化は困難であったため、ブームは下火になっていきます。

1980年代に入り、コンピュータの性能やデータ処理能力が徐々に上がっていくと、ブームが再燃。特定の問題に対して専門知識を持つ「エキスパートシステム」が登場しました。エキスパートシステムは、明確に定義できる問題に対して最適な解決策を見つける探索能力に



は長けていたものの、実社会の複雑な問題に対応できるまでには至らず、実用化には課題がありました。

2000年代半ばになると、コンピュータの情報処理能力が大幅に向上したことで機械学習が発達し、大量のデータからコンピュータが自律的に学習して特徴を判別する機械学習方法「ディープラーニング」が開発されます。ディープラーニングは、AIのプログラムに人間の脳の仕組みをシミュレートさせる「ニューラルネットワーク」を応用した技術。ディープラーニングによってAIの性能が飛躍的に進歩し、LLMなども含む生成AIが誕生したのです。

研究者も驚きの急激な進化

――急激にAIが進化したターニングポイントは何だったのでしょうか。

黒橋 最も大きな要因は、先ほど申し上げたニューラルネットワークの出現です。

AIの賢さの本質である言語を理解する能力、いわゆる「言語モデル」とは、入力された単語に対して、うしろに続く単語の確率を統計的に予想するモデルのこと。生成AIの基盤となる技術で、機械翻訳やテキスト生成などに広く活用されています。

従来は、与えられた大量のテキストデータの集合（大規模コーパス）から、直前の数単語の文脈を考慮して次の単語を予想するモデルが研究されてきました。ただし、これは各単語を記号として認識しているため、どれだけ大きな

日本でのAI発展を目指して

――黒橋さんは、2024年4月に新設された「大規模言語モデル研究開発センター」（LLMC）にてセンター長を務められています。現在、どのような活動をなさっているのでしょうか。

黒橋 2022年11月末にOpenAIがChatGPTを発表し、自然言語処理の研究者たちとの間で、なぜこんなに賢いAIができたのかという話になりました。OpenAIは技術情報を非公開にしているためです。そこで、とにかくLLMの勉強をしなくてはと、2023年5月に立ち上げたLLM勉強会が、LLMCの前身です。LLM勉強会での研究が順調に進んだこともあり、2024年4月にLLMCを発足。LLM勉強会も、「LLM-jp」という名前で活動を継続しています。

LLMCでは現在、LLM-jpと連携してLLMの開発に取り組んでいます。OpenAIのモデルがこれほど賢い理由を、実際に自分たちで作って調べてみようというわけです。

加えて今後、LLMの透明性や信頼性の向上に関する研究も進めていく予定です。ニューラルネットワークの情報処理の過程を観察するのはほぼ不可能。しかし、LLMの学習プロセスにおけるモデルの変化を追跡する、LLMの回答が学習したテキストのどの部分に基づくものかを検索によって特定する方法をとることで、LLMの回答の根拠を知ることができ、さらに、事実ではない情報を生成してしまう「ハルシネーション（幻覚）」などの現象の原因を

コーパスを与えたとしても、基準として認識できるのは直前の2〜3単語程度のみ。したがって、やはり出力テキストの精度はそれほど高いものではなく、機械特有の不自然な文章になってしまっていました。それが、ニューラルネットワークによって激変したのです。

――ニューラルネットワークが言語モデルにもたらした影響はどのようなものですか。

黒橋 ニューラルネットワークモデルで画期的だったのは、単語を単体の記号ではなく、数千、数万次元上のベクトルで表現したこと。例えば「りんご」と「みかん」のような共通項の多い単語は、近い座標値で表わされます。ベクトルという意味空間ができたことで、言葉の意味を扱えるようになったのです。

こうして、単語をベクトルで表現した初期のモデルの「Word2Vec」や、「再帰型ニューラルネットワーク（RNN）」を用いた言語モデルなどが開発されますが、いずれも長い文脈を正確に扱うことはできませんでした。こうした弱点を克服し、LLMの礎となったモデルが「Attention」であり、続く「Transformer」です。

Attentionは機械翻訳において、文章のすべての単語を同時に認識し、その名のとおり、元言語のどの単語が重要かに「注目」する能力を持っています。さらに、Attentionを精緻化したTransformerは、長く複雑な文脈を正確に処理できるようになり、ChatGPTなど、以降のLLMのほぼすべてのベースになっています。

Transformerの登場を機に、LLMは驚くほど自然な言語表現が可能となりました。最近

特定することにも繋がります。また、生成AIの安全性も、非常に重要なテーマ。LLMの理解力を高めることで、モデル自体に自身の有害性を抑える能力を持たせるアプローチも検討しています。

AIの透明性や安全性を十分に検討することは、利益を追求する企業ではなかなか難しく、多くの世界的企業がクローズドな研究環境で開発を進めています。けれども、人々の目が容易に届かないなか、秘密裏に研究が進み、ある日突然、人間がコントロールできないAIが登場するといった事態は望ましくありません。LLMCでは、すべての研究成果を開示しています。研究開発や議論の場をつくり、情報を提供することで、日本企業のAI研究の発展に貢献していきたいと考えています。

AI活用への社会受容

――AIの開発競争は今後さらに熾烈を極めていくと予想されますが、黒橋さんの思い描く理想のAI社会とはどのようなものでしょうか。

黒橋 これからの社会において、AIの普及がより勢いを増していくことは疑いようがありません。AIによる業務の効率化や高度化は、人手不足やコスト面など多くの社会課題解決のための優れた手段であり、現在、AIの導入率が低い業界でも、着実に活用が進んでいくと考えられます。

確かにAIは完璧ではありません。ハルシネーションなどの現象が問題視されているほか、AI兵器といった軍事利用への懸念もあり

では、コーパス量やニューラルネットワークのパラメーター数（モデルのサイズ）もどんどん増えて、この2〜3年で急速にその精度は上がっています。こうした進化は、LLMの研究者の誰もが目指していたことですが、実際にここまで急に賢くなるとは誰も思っていないでした。LLMが実行している内部処理自体は非常に明白であるにもかかわらず、なぜ入力に対してこうした自然な出力が可能なのか、なぜこれほど急に賢くなったのか、実は誰にも明確にはわかっていないのです。



ます。しかし、こうした問題を回避するために、AI開発を遅らせることが正しいかといえば、そうではないと思います。また、一方で、パラメーター数やコーパス量など、物量的なAIの進化は、現状ではやや頭打ち感があります。きちんとした技術力を保持しつつ、人間とコミュニケーションが可能」というAIの賢さを、適切な形で社会に浸透させていく方法を模索し続けることが必要なのではないでしょうか。

――実社会への活用もさらに進んでいくと思いますが、企業や団体がAIを導入する上での課題や、求められる対応について教えてください。

黒橋 画一的にこうしなければならない、というようなものではなくて、それぞれのシーンごとに、副作用なども考慮して判断していくしかないと思います。加えて、社会が受容するかどうかが問題です。

例えば、車の自動運転などはすでに、人間による運転よりも事故率が低いことが受け入れられて、実用化され始めていますよね。一方で、教育業界においては、数学や英語など、AIなら教師と同じレベルの指導が可能ですが、心理的な拒否感があるためか、現時点での活用範囲は限定的です。医療業界など人命に関わる分野では、さらに心理的ハードルは高くなるでしょう。確かに、AIにもミスはあります。けれど

も確率的に言えば、可能性はかなり低い。大切なのは、どの範囲までA Iに任せるかを法律等できちんと決めることだと思います。

かつて防犯カメラの導入の際にも、最初はプライバシーの問題でかなり反発があったといいます。しかし、犯罪抑止に大きく貢献し、その実績によって防犯カメラがある社会が普通になっていきました。これはなにも、防犯カメラ自体の精度が上がったためではありません。社会の皆さんの考え方が徐々に変化し、受け入れられていったのです。

A Iを活用するメリットは非常に大きいと思

います。運用上のルールなども整備しつつ、ぜひうまく活用していただきたいです。

――首都高速道路をはじめとするインフラ業界に期待することはありますか。

黒橋 構造物の劣化や渋滞予測の支援、交通状況を監視するための活用などの面で、A Iの活用はますます進んでいくのではないかと思います。先ほども申し上げたとおり、A Iにも間違いはあるため、危険度によって活用の範囲を検討する必要がありますが、どこでどのように使っていくかを積極的に考えていくことが重要だと思います。



くろはし・さだお／京都大学大学院工学研究科博士後期課程修了後、2006年、京都大学大学院情報学研究科教授に就任。2023年に国立情報学研究所の所長に着任したのち、2024年4月からは大規模言語モデル研究開発センターのセンター長も務め、自然言語処理やLLMの研究開発を行っている。著書に「自然言語処理」（放送大学教育振興会）。