

講師 **松原 仁氏**

(京都橘大学工学部情報工学科 教授)

演題 **人間の概念を変える
人工知能**

令和7年4月17日(木) 開催



はじめに

こんにちは、ご紹介いただきました松原です。

本日は“AI（人工知能）”の最近の話題をお話したいと思います。タイトルは、最近の進歩した人工知能が人間の概念を変えつつあるという見方を踏まえて、「人間の概念を変える人工知能」としました。よろしく願いいたします。

人工知能（AI）とは何か

本日のテーマである「人工知能」は、Artificial Intelligenceという英語の日本語訳です。最近では日本でも頭文字を取って“AI”と言われることが多くなったと思います。

最初にお断りしないといけないのは、「AIとは何か」ということが、専門家の間でもコンセンサスがないということです。私がAIを研究し始めた40数年前、大学院の時には先輩たちに「ない」と言われましたが、未だにありません。「知能を人工的に実現する」という目標はかなりはっきりしていますが、「知能とは何か」ということは、辞書を読んでみても分かったような、分からないような説明なのです。

知能にはいろいろな側面があります。例えば、言葉を話す、読んで理解するという側面です。コンピューター言語である「人工言語」に対し、我々が使っている日本語や英語を「自然言語」と呼んでおり、これをAIの分野では「自然言語処理」と言います。また、資料を目で見て理解するという「画像認識」、耳から入ってきた情報を理解、認識するという「音声認識」

があります。

今申し上げただけでも3つの機能を知能と言っており、これ以外にも、人とうまくコミュニケーションするとか、お金を数えるとか、いろいろなことが知能であると言われております。何ができることが「知能がある」ということなのか、たくさんありすぎて数え上げることが事実上できないのです。

私は幼稚園の時に漫画「鉄腕アトム」を見て、将来はこんなロボットをつくる仕事がやりたいと思って結果的にAIの研究者になりましたが、何ができるか「人工知能ができた」と認めてくれるのか、よく分かっていないのです。

鉄腕アトムみたいな、誰が見ても、誰が付き合っても、このロボットは人間並みの知能があるというロボットができた暁には、そのロボットの能力が知能だと言えます。知能は人間の中にあるから客観的に取り出すことはできないのですが、もし人間の外に知能があれば、ようやく客観的に知能について議論できるのではないかと、そういう夢や野望をAI研究者は持っています。

AI研究は「知能を持った人工物」を作るという工学的な目的と、その過程を通じて知能とは何かを理解するという科学的な目的の両方を持っているのです。

今、AIがブームなので、AIと言った方が売れると思えば、これまでアプリと言っていたものをAIと言う傾向はあります。それをあまり責められないのは、そもそもAIというものの定義が明確ではないから、という面があると思います。

80年代にワープロで、かな漢字変換が出てきた時は最先端のAIと言われていましたが、今となっては、あ

んなのはできて当たり前で、かな漢字変換をAIと言う人はもう誰もいなくなっています。それはAIの技術が進歩したためです。

そういうことからすると、今AIと言われてもてはやされているものも、将来AIと言われなくなる可能性はあると思います。

人工知能（AI）の歴史

1. 1940年代 コンピューターの発明

人工知能（AI）の歴史は、ほかの学問分野と違って、非常に短いです。AIはコンピューターを使っていますので、AIの歴史はコンピューター発明以降ということになります。

1940年代、第二次世界大戦の末期に、イギリスやアメリカが主に戦争目的で、例えば、ミサイルを動く目標に命中させるには何度か角度でどれくらいの速度であればよいのかなど、数値計算を早く行うために、コンピューターが発明されました。

2. 人工知能研究のスタートと1回目のAIブーム（1950年代～1960年代）

人間は内省すると、頭の中で概念や言葉を「ああでもない、こうでもない」と考えているという印象があります。そこで、コンピューターは数値計算だけでなく、記号も扱えるのではないかと、例えば言葉や概念のようなものをコンピューターで扱えるなら、人間のよう賢いことが考えられるのではないかと、1950年頃にイギリスやアメリカで新しい分野として研究が始まりました。

最初の頃は名前がなかったのですが、1956年にジョン・マッカーシーというアメリカのAI研究のパイオニアの一人が「この分野をArtificial Intelligence（AI）と呼ぶことにしよう」と提唱し、この呼び方がシンプルでキャッチーだったということもあって、今も使われています。

最初のAIはアメリカとヨーロッパにおいて、すごい予算をつけて研究者を集めて、盛んに研究されました。

しかし、当時はコンピューターの能力を過大評価しておりました。当時のコンピューターは能力が低く、世界最高のコンピューターでも皆さんがお持ちのスマー

トフォン1台よりもはるかに能力が低いものでした。

それと同時に、人間の能力を過小評価していました。先ほど知能の例として自然言語処理、画像認識、音声認識と申し上げました。例えば、日本語を話すとか、聞いて分かるというのは、日本語を母国語にする人にとってはほぼ苦勞がありません。難しい話ならともかく、世間話であれば、自分自身、何で分かっているのかもよく分からないけれど、分かります。聞き流していても分かるのです。画像認識もそうです。ペットボトルを見て、なぜペットボトルと分かるのか、そう聞かれても説明のしようがないですね。だって、ペットボトルでしょう、みたいなことになります。AIのテーマの多くは「人間がなぜそれをうまくできているかが分かっていないこと」なのです。

人間にとっては無造作にできることでも、コンピューターにとっては「どれも非常に難しい」ということが1回目のブームの時に分かりました。それが成果といえば成果なのですが、役に立つものはほとんどできなかったという意味ではやっぱり失敗なので、研究費を削られ、研究者も少なくなり、1回目のブームの後、すごく厳しい「冬の時代」となりました。

3. 2回目のAIブーム（1980年代～1990年代）

1980年代になって2回目のブームが来ます。この時は“エキスパートシステム”という専門家の代わりをするシステム、特に医療診断におけるシステムが注目されました。内科の患者のデータを入力すると、病名を推定し、適切な薬を処方するシステムでは、その分野の専門の先生よりは劣るが、若い先生よりは成績が良いという結果が出たのです。

そこで医療だけではなく、法律や金融などいろいろな分野にこのシステムを活用しようという動きが世界中で盛り上がりました。この頃は、日本の景気がよかったのが日本が一番研究費をかけたのですが、残念だったのは、当時のコンピューターの能力が今のようないレベルに達していなかったということです。このときうまくいけば、AIは日本の天下だったと思うのですが、優秀な専門家のレベルに到達できないのでは使い物にならないということで、また冬の時代に入ります。

4. 3回目のAIブーム（2010年代～）

2回目の冬の時代の最中の2006年に、“ディープラーニング”という、今のAIブームの基礎になっている技術をジェフリー・ヒントンのカナダのトロント大学の先生が開発しました。我々専門家の間ではすぐ話題になったのですが、一般に使われるようになったのは2010年代です。すごくいろいろなことができるというので、3回目のブームになりました。

3回目のブームから10年経ち、そろそろ落ち着くかと思いきや、生成AIが2022年に出てきて、また大騒ぎになり、現在に至っております。

ディープラーニングとは

1. 2006年にヒントンが提唱

2024年にノーベル物理学賞を受賞したことでジェフリー・ヒントンの名前はご存じかと思います。

彼は1980年代の2回目のAIブームの時から研究している人で、その時にニューラルネットワークという当時の技術としては画期的な、今のディープラーニングの原型のようなものを作りました。そのことがノーベル物理学賞の受賞理由に書かれていますが、AI研究者は皆「今のディープラーニングの隆盛に貢献したこと」が実質的な受賞理由だと思っています。

ディープラーニングは2006年にヒントンがゼロから思いついたというよりは、1950年代、AIが始まった時からその源流がありました。人間には数百億個の神経細胞があると言われており、シナプスという線で結ばれて、複雑なネットワークを頭の中に形成しています。それが知能の源泉なわけです。それなら、そういうネットワークをコンピューター上にシミュレーションすれば、コンピューターは賢くなるはずなんです。これは別にすごいアイデアではなく、誰でも思いつくものですので、AIが始まった頃から既にある考え方です。

でもコンピューターの能力が貧弱だったので、すごく簡単なシミュレーションしかできず、最初はパーセプトロンであり、少しコンピューターの性能が良くなってから開発されたのがニューラルネットワークでした。今回のディープラーニングはこれらの拡張版です。パーセプトロンやニューラルネットワークに比べ

て、シミュレーションのネットワークの層が厚くなったという意味で「ディープ」なのです。

2. ディープラーニングの特徴・問題点

ディープラーニングは、人の顔の認識とか、人が話したことの認識とか、パターン認識が得意です。今は、自然言語処理も扱います。

「ディープラーニングは何をやっているのか」というと、簡単に言うと、たくさんのデータから傾向を適切に見いだすことができるのです。だからディープラーニングは優秀ですが、データがたくさん必要です。

ディープラーニングは「パフォーマンスは良いけれど、なぜその答えになったのかが分からない」とよく言われます。すごく複雑なネットワークなので、人間が解析できないわけです。例えば、囲碁でいい手を打ちますが、何がどう計算されて、その手になったかが分からないのです。

これは大きな欠点なので、説明可能AIという分野が世界中で盛んに研究されており、ある程度は説明できるようになってきています。

最近のAI技術の進展

1. 個人認証

人の顔の認識が人間の精度を超えつつあります。今、税関においてもパスポートと顔が一致しているかをまずAIが見て、AIが怪しいと判断すると、税関職員に伝えて、別室に連れていかれる、そういう風になっているようです。

2. 会議の（音声）記録

オンライン会議などの文字起こしもかなりの精度でできるようになっています。

翻訳も同時にできるようになっており、外国人が参加している会議で、日本人が話したことを英語で画面の下に出し、外国人が話したことを日本語で画面の下に出すことが、だいたいリアルタイムでできるようになっています。

3. 自動運転

自動運転についても、アメリカや中国では自動運転

タクシーが実用化されております。乗った人の話では、ブレーキを踏むタイミングが人間と違うため、最初はちょっと怖いとのこと。人間は早めにブレーキを踏みがちですが、AIは安全に停止できるぎりぎりまで踏まないということがあるようです。

4. 株取引

株取引についても証券会社がAIを使っております。過去のデータを大量にAIに入れて、過去の同じような状況の時に、どこの株が上がった、下がったということを見て売買を判断しているようです。

5. 創薬・化学

創薬にはすごくお金がかかるので、AIでシミュレーションして、コストを下げっております。

化学ではAlphaFoldというシステムが注目されています。専門家に言わせると、タンパク質の立体構造が薬効に影響があるということですが、このシステムではすごく早く正確にタンパク質の構造予測を行います。この研究が2024年にノーベル化学賞を受賞しました。

人工知能が得意なこと、苦手なこと

1. 定型的・標準的作業は得意

先ほど申し上げたように、AIはたくさんのデータから学習しているので、データが多い領域では強いのです。だから定型的な、標準的な作業は得意です。

標準レベルというのは、実例が一番多いので学習しやすいわけです。一方、トップレベルの人のやっていることを学習させようとなると、データが少ないので、AIでは一筋縄ではいかなくなります。

2. “意味”は理解できない

今の生成AIもあれだけ流暢に答えていますが、文章の意味は全く分かってないのです。我々は生成AIがどういうアルゴリズムで動いているか知っているので、断言できます。意味が分からないのに、あれだけ流暢に意味が通じているようなことを言うことが、あの技術のすごいところです。

3. ルールが明確なこと、範囲が限定されていることは得意

AIはルールが明確なものが得意です。だからゲームに強いのです。範囲が限定されているものも強いのです。

でも、世の中の多くの問題はルールが不明確で範囲が非限定です。AIの最先端はルールが不明確、範囲が非限定なものをうまく解こうとしているのですが、人間の専門家にまだ劣るところが多いと思います。

4. 理性的なことは得意

AIは理性的なことが得意です。感性的なことはまだ苦手なのですが、生成AIなどは今、小説モドキを書き始めているし、絵も画像も描いているので、そろそろ感性的なこともできるようになりつつあります。

AIで蘇る手塚治虫（Tezuka2020とTEZUKA2023）

今から5年前になりますが、「Tezuka2020」というプロジェクトに参画しました。手塚治虫さんがもし生きていたらどんな漫画を描くだろうか、それを現代に蘇らせようとするプロジェクトです。手塚プロと一緒に参加したので、手塚さんの作品に関する著作権の問題はないということで、鉄腕アトムやいろいろな作品のシナリオを入れて新作の候補を生成しました。「ばいどん」という漫画になり、発表されたのですが、生成AIはまだなかったので、AIの果たした役割は1割程度にとどまりました。

2023年の経済産業省関係の研究プロジェクト「TEZUKA2023」では、手塚治虫さんの「ブラックジャック」の新作を作ることになり、これも手塚プロと組みます。これは生成AIを利用したので、3割ぐらいはAIが貢献できたかなと思っています。

言語生成AI、ChatGPTの登場

2022年11月にChatGPTがリリース（公開）されました。言語生成AIの代表格です。OpenAIというアメリカの会社が開発したものです。

GPTというのは技術の名前です。Gは“generative”すなわち「生成」、Pは“pretrained”すなわち「事

前学習」、最後のTは“transformer”これは手法の名前です。適切な訳がないので、日本でも「トランスフォーマー」と言われることが多いのですが、ディープラーニングの一種です。これが肝なのです。

OpenAI社は公開されているネット上の文章を大量に集めます。何を集めたかは公表しておりません。著作権者からの許諾は取っていないし、使用料も払っていないのです。多分アメリカの有力新聞の記事は絶対使っているはずですが。だから今、アメリカの多くの有力新聞社がOpenAI社を裁判で訴えており、裁判の結果が注目されております。

OpenAI社がお金持ちだと思うのは、世界中で何億人もの人々がChatGPTを使っているにもかかわらず、無償で使わせていることです。すごいことだと思います。

ChatGPTは無償ですが、その改良版「GPT-4」は月20ドル、最先端の「ChatGPT o1（オーワン）」は月200ドル必要です。高いものほど良い答えを出しますが、無償のものでも今では良い答えを出します。

言語生成AIの仕組み

生成AIの基本的な仕組みは、実は単純なのです。文章の中で一単語を外して、「そこにどの単語が来る確率が一番高いのか」を計算して、それを連続的に行って文末までいきます。

我々AI研究者もトランスフォーマーというアルゴリズムは常識で知っていたのですけれど、それを使ってあのような流暢な言葉が出てくるとは誰も信じていなかったのです。

やはり肝はデータの量です。覚えさせた文章がすごい量だったということです。「日本の総理大臣は」の次に何が来るか、となると、「日本の総理大臣は石破さんだ」と書いてあるのが新聞記事等にたくさんあるわけですね。その確率が圧倒的に多いので、ほとんどの場合「石破茂」というのを持ってきます。「総理大臣岸田文雄」と書いてある記事等も学習した中に一定程度残っているとすると、ある確率で「岸田文雄」と書いてしまう。これが「AIは時々間違ふ」と言われる主な理由です。ずいぶん工夫して間違いは減ってきましたが、統計的に処理している以上、間違ふ可能

性をゼロにはできないのです。

また、変なことを書かないように学習させています。個人のブログなどからも学習するので、差別的なことや猥褻なことを書いている文章がそのまま出くると困りますから、そういうことは言わないように学習させているのです。

ChatGPTの特徴

1. 日本語の入出力が可能

ChatGPTは日本語の入出力が可能ですが、処理は英語（日本語を英語に翻訳して処理して、結果の英語を日本語に翻訳して出力する）です。

今、日本版の生成AIもいろいろ出てきております。もともと日本語の文章を入力しているので、欧米製の生成AIよりも日本語が流暢だったり、日本のことに詳しくたりします。

2. 日本についての知識は少ない

ChatGPTは最初の頃は日本のことを尋ねると、結構間違えておりました。データに日本のことが少なかったということですが、今はChatGPTにも昔よりは日本のデータが入っています。

3. プログラムや長い文章も書ける

ChatGPTはプログラムも書けます。私も情報系の大学の教員なので、学生に「こういう入力をして、こういう出力をするプログラムを2週間後までに書いて提出しなさい。」と宿題を出すのですけれど、今では生成AIが学生に代わって完璧にそうしたプログラムを書いてしまいます。それくらいのレベルにはもう来ていて、かなり長い文章も書けるようになっております。

4. とときどき嘘を言う

「ときどき嘘を言う」とは、間違えるということであり、これが一番の欠点だと言われております。これは先ほどご説明したように、統計的に処理している以上はゼロにはできないのですけれど、だいぶ減ってきております。心配な場合は、同じ質問を複数の生成AIに聞いて、みんな答えが同じであれば、だいたい信じていいだろうし、生成AIによって違うことを言

うようなら、ちょっと考えてみる、という対応をすれば十分に使い物になると思います。

5. 要約や翻訳も得意

要約や翻訳も得意です。990点満点のTOEICでは、3年前ぐらいでAIは950点を超えたと言われているので、今ではほとんど満点に近いと思います。私も外国人と英語のメールのやり取りがあるのですが、ちょっと微妙な表現については、生成AIの英語を使っ、最後は自分でチェックする、ということもやっております。

論文についても、日本語で論文を書いて生成AIに訳させてみると、だいたい言いたいことが英語で書かれているのです。

英語が苦手な日本人としては、生産性を上げるという目的からすれば、よい道具だと思います。

GPT-4等の能力

ChatGPTの改良版GPT-4は、先ほど申し上げたように月20ドル支払う必要があるものですが、登場してすぐ、2023年頃にアメリカの司法試験、医師国家試験に合格しています。日本の医師国家試験にも合格しています。

日本の司法試験に合格できないのは、やはり欧米の生成AIなので、日本の法律に関する知見があまりないからだと思います。今、日本の生成AIが日本の司法試験合格を目指して頑張っているとのこと。

さらに、OpenAI社の最先端の「o1（オーワン）」及び「DeepSeek」という中国製でちょっと話題になっている生成AI、この2つを使って、今年の東大の二次試験、理科一類から文科三類まで全部を解かせて採点したら、全部合格点に達したとのこと。こういうペーパーテスト系に関しては、今やAIは十分合格点レベルに達しているのです。

生成AIとどう付き合うか

私たちは発展目覚ましい生成AIとどう付き合っていくたらよいか、ということについては、私自身は自動車の比喻で考えるのがよいと思っております。

それまではみんな馬車で移動していたのですが、世界最初の自動車であるT型フォードが登場して、すごく便利だったため、いろいろな会社が瞬く間に真似をして自動車を作り始め、現在の自動車社会に至っております。規制しようが何をしようが、作ることができると分かってしまうと、作り方は教えてくれなくても、技術というものはできてしまうのです。

昔、出来上がった鉄砲を入手した日本の技術者が、見よう見まねで鉄砲が作れるようになりました。生成AIについても、出来上がった生成AIがあれば、作り方を教えてくれなくてもできるようになるので、禁止は意味がないと思います。

自動車については、免許もないし、制限速度もない、そういうところからルールを作り、今の自動車社会ができました。今でも死亡事故が一定数ある危険な機械ではあるのですが、メリットとデメリットを考えて、自動車を道具として使ってきたのです。

生成AIについても、統計的な処理という側面や悪用される恐れもあるわけですが、とても便利な道具で、生産性を非常に上げるので、好むと好まざるに関わらず、道具として使っていくということかと思っています。

DeepSeekとは

先ほど少し言及したDeepSeekは、中国の会社が作った生成AIで、2024年の冬に発表されたものです。DeepSeekはオープンソースであるという特徴があり、これは他人がプログラムを読めるということです。ChatGPTでも公開していないのに、中国の会社が公開するのは画期的です。

技術的なポイントは「強化学習」と「蒸留」です。蒸留という技術自体は生成AIを勉強している人なら誰でも知っているものですが、DeepSeekはそれをうまく使って、高性能の生成AIをとて安価で作ったことが話題になっています。

データが多ければ多いほど性能が良くなるという経験則が生成AIにはあります。だからOpenAI社はお金を投資家からたくさん集めて、データを増やして性能を良くしています。たくさんデータを集めると、高性能のコンピューターが必要となります。そのた

め、その分野の最先端である NVIDIA という会社から高価な GPU (Graphics Processing Unit) を OpenAI 社が買い上げるという構図ができているのです。

DeepSeekの発表直後にNVIDIAの株価が大暴落したのは、「NVIDIAのGPUがなくても高性能な生成AIを作ることができるじゃないか」ということが原因です。

生成AIがすごく安価で作れるようになったことは画期的なことだと思いますし、お金のない組織にもチャンスがあるということです。今、DeepSeekのやり方で日本も追従しようとしています。

今後の生成AI

今では言葉だけではなく、画像や音声も生成AIで扱えます。OpenAI社のo1などはそうです。

もともと生成AIは、何を聞かれてもそれなりに答える万能型だったのですけれど、今では医療や法律など分野を限ったものを作ろうという話があります。質の良いデータ、医療の正しい知識のデータだけ集めればいいし、医療だけだったら開発にそれほどお金もかからない、著作権についても文句を言われる前にデータ提供者に許可を取ってしまおう、という流れになっております。

大きい会社や自治体、国などの組織では、情報漏洩が問題になることがあります。プロンプトに企業秘密を入力すると、それが学習に使われて、回りまわって社外の人に使われてしまうのではないかと危惧されるのです。それを使われないようにする技術的なテクニックもあるので、外のネットワークから離して、組織内だけでその生成AIを使えばよいのです。組織の一番大事な情報を入れることで一番詳しくなってくれます。そうすれば組織として非常に使いやすい生成AIができて、外に出さない限り秘密も守れる、ということになるのです。OpenAI社がやっているような巨大型に対して、コンパクト型と呼ばれています。

OpenAI社のスポンサーにMicrosoftがついているので、遅かれ早かれ、WordやExcelに生成AIが標準装備されると思います。今後はExcelでもWordでも生成AIから書き方を提案されるようになることが普通になり、本人からすると、生成AIを使っているか否かを意識することがなくなってくると思います。

将棋における人間とAIの関係

私は将棋のAIの研究を、AIがとても弱い時からやってきました。今はご存知のとおりAIの方が強くなりましたが、AIが人間の實力に追いついた頃の2000年代後半から2010年代半ばぐらいまでの間は、かなり軋轢がありました。「プロ棋士がコンピューターごときに負けることがあればプロ棋士の尊厳に関わる」、「AIが勝つようになるとプロ組織は存続できるのだろうか」などと、プロ棋士や将棋ファンがAIに対して反感を抱くようになったのです。

しかしAIの實力がプロ棋士を超えたので、軋轢はなくなりました。今ではプロ棋士や将棋ファンはAIで勉強しています。とても良い関係にあると思います。

例としてプロ棋士の藤井聡太さんと永瀬拓矢さんのタイトル戦をご紹介します。これが結構面白いなと思ったのは、お互いAIで予習してから対局に臨んでいるのです。

75手目までは定跡通りです。このままいくと、先手の藤井さんが勝つということは、AIがほぼ解明してしまっています。

後手の永瀬さんはどこかで変えないといけない。そこで永瀬さんは定跡を外れる九五歩を打ったのです。永瀬さんとしては「藤井さん、この先はあなた読んでないでしょ。私はこの先を勉強したのですよ」という気持ちです。

さらに80手目で永瀬さんは八四桂馬を打って、銀取りになっています。銀を取るとそれが王手になっているので、藤井さんは一見ピンチです。ほとんどの人は銀をどこかに逃がさないといけないと考えます。ところが藤井さんは全然逃げずに四六香と打って攻め、藤井さんが勝つのです。

永瀬さんはすごくショックだったと思います。自分が途中で間違えたならともかく、思ったとおりに進んで、気が付いてみたら負けの局面だったのです。

AIには選べても、人間なら怖くて選べない将棋の指し手があるのですが、藤井さんはそうした「肉を切らせて骨を断つ」という将棋の指し方をして、すれすれのところで勝つのです。だから強いのだと思います。

今、藤井さんを含めてプロ棋士は、皆さんAIで勉強しており、10年前のプロ棋士より強くなっていま

す。ファンから見てもレベルの高い将棋が見られるので楽しいと思います。

これからの人間とAIの関係

まず、道具としてのAIが人間を助けるということ、次にAIによって人間も変わるということです。

将棋の例でお分かりのように、賢い道具としてAIを使うことで、人間がさらに成長できることを示していると思いますし、だからAIは社会の多様性確保のための要素の一つなのです。

一方で、将棋界とAIとの間にかつて軋轢があったように、これからいろいろな分野でAIが進歩していくと、AIを受け入れたくないという軋轢が生まれてくることでしょう。人間の本能として、自分が大事だと思っているもの、自分が仕事の材料にしているものが、自分以外のものにもできるようになるというのは、尊厳を冒される、存在が脅かされるという気持ちになるのは当然なのです。今、生成AI関係でいくつもSNSが炎上していますけれど、過渡期でやむを得ないかなと思う側面もあります。

だから人間が新しいものを受け入れるには、ちょっと時間が必要なのかなと思います。「広い心でAIとも付き合おう」とはそういう意味なのです。

人間という概念

スマートフォンを持っている人間は、電話番号にしても、漢字を書くにしても、スマートフォンに結構情報を任せていますよね。また、出張時においしいものを食べようと、関連するサイトをつい覗いてしまう、というように意思決定のかなりの部分をスマートフォンに任せているという実情もあります。

確かに、人間がますます思考しなくなるんじゃないとか、いろいろと問題はあって、考えなければいけないことは山積しています。

でも、それがないよりは能力は上がっているはずなので、それでいい、それを良しとしようと考えて、人間という概念をアップデートしてみる。単独としての人間から、AI機器を含めた存在としての人間になる、すなわち、将来的には「人間＋スマホ（の延長線上の

ようなもの）」が新たな人間という概念になっていくのではないかと、世の中の流れとしてはそういう方向に進むのではないかと、思っております。

ご清聴ありがとうございました。



講師略歴

松原 仁（まつばら ひとし）

京都橘大学工学部情報工学科 教授

1981年 東京大学理学部情報科学科卒業

1986年 同大学院工学系研究科情報工学専攻博士課程修了 工学博士

1986年 通産省工業技術院電子技術総合研究所（現産業技術総合研究所）入所

2000年 公立はこだて未来大学システム情報科学部教授

2020年 東京大学次世代知能科学研究センター教授

2024年 京都橘大学工学部情報工学科教授 現在に至る

元人工知能学会会長、元観光情報学会会長、前情報処理学会副会長
著書に「鉄腕アトムは実現できるか」、「AIに心は宿るのか」など