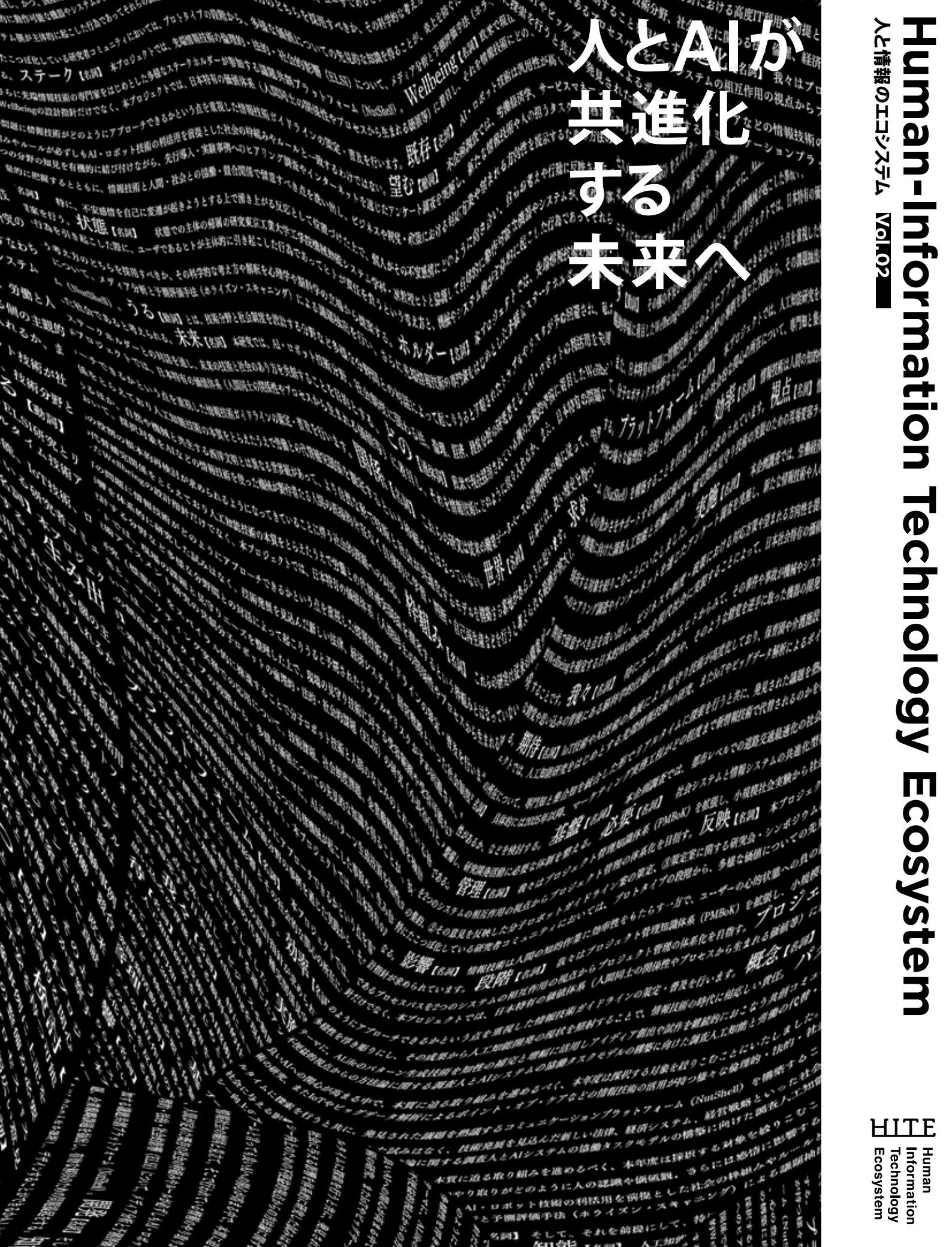


人とAIが 共進化 する 未来へ



領域総括挨拶

AIの力を過大に評価して、その社会インパクトを大きさに語るのは恐らく非科学的なことなのだろうと思います。生産性向上や自動化による雇用破壊などは、ラッダイト運動を持ち出すまでもなくAI以前から存在していましたし、それに対応した社会制度を作るなどして、経済的に豊かな社会を作ってきました。また、人間はすでにさまざまな自動制御システムに判断を委ねながら生活をしています。AI固有の新しい問題というのは案外少ないのかもしれない。

そんな冷静さを保ちつつも、AIに代表される今日の情報技術が、「判断のブラックボックス性を高めて人間にとって制御不能になるのではないか」、「悪意ある利用のされ方をした時に人類に脅威を与えるのではないか」、「中間層の雇用破壊や、システム（知財を含む）所有を軸に、貧富の格差を拡大させるのではないか」、「プロファイリング技術が進んで相互監視社会化が進み人権が侵害されるのではないか」などの懸念は根拠がないとは言えません。AIの可能性が大きい分、それを人類にとってより良い結果をもたらすものとして、育て、受け入れていく知恵を持ちたいものです。

対応として真っ先に思いつくのは、技術開発に対して何らかのガイドラインを設けることですが、それは簡単なことではありません。技術進化の方向が完全には予見できないからです。また、技術を受容する側の社会も思わぬ展開をすることが歴史的に知られています。ルールの中核となる人間の価値観も変化していきます。ビジネス化がどのように行われるかというのも、技術、社会両面に大きな影響を持ちます。そんな中で、RISTEX「人と情報のエコシステム」領域では、技術も、社会制度も、人も共に進化するものととらえて、理解と対話を通じて人と情報技術が共生できる（なじむ）、プラットフォームづくりを目指します。

その目標に向けて、基盤となる概念構築、分野をまたがった研究者のコミュニティづくり、ステークホルダーが相互にフィードバックできるプラットフォームづくりなどを目指して活動を行っています。

領域総括

國領 二郎



國領 二郎

「人と情報のエコシステム」領域総括
慶應義塾大学 総合政策学部 教授

1982年東京大学経済学部卒業。日本電信電話公社入社。92年、ハーバード・ビジネス・スクール経営学博士。2003年から慶應義塾大学環境情報学部教授などを経て、09年総合政策学部長。05～09年まで慶應義塾大学湘南藤沢キャンパス研究所長も務める。2013年より慶應義塾常任理事に就任。主な著書に「オープン・アーキテ クチャ戦略」（ダイヤモンド社、1999）、「ソーシャルな資本主義」（日本経済新聞社、2013）がある。

Human-Information Technology Ecosystem

「人と情報のエコシステム」は、
科学技術振興機構（JST）社会技術研究開発センター（RISTEX）が推進する
研究開発領域であり、今年で3年目を迎えます。AI、ロボット、IoTといった情報技術が
加速度的に進展する現在、いかにそれら技術が社会に浸透し、人間の暮らしになじんでいくか、
またそのときどんな問題が起きうるかを考え、新たな制度や技術を設計していくための
研究開発を推進します。情報技術と人間が、よりよい関係を結んでいく社会を
実現することが、本領域の目的です。

Website: <http://ristex.jst.go.jp/hite/>



特別鼎談

人工知能時代の責任と主体とは？

松浦 和也

哲学

よしだ 貴子

心理学

稲谷 龍彦

法学

もし、人工知能が事故を起こしたら、誰が責任を取るのでしょうか？

これは情報技術が急速に社会に浸透するいま、未だ決着のつかない重要な課題となっています。

この現代的問題に挑むべく、哲学、心理学、法学と異なる専門性を持つ3人の研究者が集い、

人工知能時代における「責任」と「主体」とは何かが議論されました。

— まずは、皆さんの研究内容と「人と情報のエコシステム（以下、HITE）」で推進されているプロジェクト内容についてお聞かせください。

松浦 和也(以下、松浦): 私の専門はアリストテレスを主とするギリシアの哲学です。ギリシア哲学という大昔のここのように感じる人もいますが、私は現在の科学観にも通底する当時の考え方がテキストの中に散りばめられていて、その考え方や価値観は今の社会をとらえ直すヒントにもなると信じて研究を進めています。そこでHITEでは、昨年度は「高度情報社会における責任概念の策定」という企画調査を、本年度からは「自律機械と市民をつなぐ責任概念の策定」と題したプロジェクトを立ち上げ、人文系の視点から「責任」の概念について考察するプロジェクトを展開しています。

このプロジェクトの特徴は、人間の歴史や文化を重視することにあります。というのも、AIをはじめとする何かしらの“自律的な機械”が社会に普及していく時、機械が実際にはどのような構造で動いており、実際にどう振る舞うかというよりも、それらを一般の人たちがいかに受け取るのかが問題になってくるからです。そこで私たちは、「自律的な機械とは何か」を定義するのではなく、「機械が人間と対等な存在であると認められるには、どんな要件が必要になるか」という問いを立てました。この問いはすなわち、「人間とは何か？」という普遍的な問いにもつながってきます。古代ギリシアや日本の江戸時代、さらには古代インドといった近代西洋以外の社会モデルも視野に入れて、人間と自律機械のあり方を見据えることによって、今後の社会のあるべき姿を提示できるのではないかと考えています。

蔵田 貴子(以下、蔵田): 私たちの研究室は、人と一緒に動作する機械やシステムの開発をしています。HITEでは、「人間とシステムが心理的になじんだ状態ではどちらに主体があるか」をテーマに、機械システムに対する人間の心理状態を観察し、機械と人の親和性に関する研究を脳科学のアプローチから進めています。例えば、身体に装着

するサポートロボットなどを使うときに、ある条件を満たすと人は次第に機械が自分の体の一部のような感覚を覚えて、自分自身が自分の身体を操作しているという感覚が強くなり、実のところ機械にも自分の身体を操作されているかもしれないという感覚が薄くなってきます。その状態で何らかの事件や事故が起こったとき、その行為は「誰に責任があるのか」という問題が浮かび上がってきます。本人からすると、すべての行動は自分自身が主体をもって行動した」感じがしてしまい、もしそれが機械側の問題であっても気付かず責任を感じてしまうかもしれない。逆に、自分が引き起こした行動の結果にも関わらず、機械の誤動作だと申告してメーカー側に責任転嫁してしまう場合もあります。本人が本当にそう思っているのか、ただ嘘をついているだけなのかを誰がどう判別するのか。さらには、自分自身の身体の行動を機械に乗っ取られたと言い張る人も現れるかもしれませんが、それは第三者からみると、機械と人が一緒に動作しているので、一見機械が人を操作しているのか、人が機械を主体的に操作しているのか判別が難しい。そうした機械と人の境界が曖昧になる状況下で、誰が特定の行為を実行したかという「主体」と「責任」をどう捉えていくかがとても重要なテーマだと考えています。



稲谷 龍彦(以下、稲谷): 私の専門は刑事法です。具体的には、哲学や経済学、社会学、認知科学といった法学と隣接する領域の知見を活かしながら、刑事法の解釈論や立法論に対して提言を行う研究を行っています。現代は、グローバル化や技術の進歩に伴って、これまで法学の世界で確固たる前提とされてきた、人間は「自由意志」を持つ理性的な「主体」であるという思想に揺らぎが生じてきています。そこで現代哲学や認知心理

学、行動経済学等の知見を生かしながら、従来の刑事司法制度のあり方を批判的に検証し、新しい刑事司法制度のあり方について考究しています。HITEで参加した「自律性の検討に基づくなじみ社会における人工知能の法的電子人格」では、最近法学の世界でも人工知能へ「法人格」を付与するべきかが議論されるようになっている現状を踏まえ、人工知能の開発・利用をめぐる法的責任のあり方、とりわけ人工知能と人間とが調和して存在するために必要な処罰制度のあり方に関する議論を展開しています。

人間に

「自由意志」はない？

— 人工知能時代の「責任と主体」について、いま最も課題となっているポイントは何でしょうか。

稲谷: 先ほども少しお話したように、西洋の近代哲学を基盤とする近代刑事法では、人間が自由意志をもって、外的環境の影響を受けずに客体をコントロールできることが基本的な前提とされています。ところが、ディープラーニングのような学習型の人工知能は、継続的に発展していくものであり、完全なコントロールは想定できません。また一方で近年の脳神経科学によれば、人間は常に外的環境の影響下にあるため、そもそも確固たる自由意志の存在が疑わしいことが明らかになりつつあります。そうすると、「何か危険が現実化した場合は、その危険源をコントロールできた人間が責任を取るべきである」という議論の前提が揺らいでいることができるのです。

蔵田: 同じく認知科学や脳科学の分野でも、私たち人間は本来、自分で思うほど自分のことを主体的にコントロールできていない可能性が議論されています。そもそも、本人にとっての主体的な感覚が、実際に起きている物理的現象と対応していないことは人間にとっては多々あります。また最近では、他人や機械にやらされた、あるいは機械が勝手にやったと感じたことについては、自分に責任がないと主張しがちな可能性も議論されています

ね。「誰に責任があるのか」という社会的な議論をする前に、こうした人間の性質をよく考察する必要があると思います。

松浦： 責任や主体という概念自体が、まさに近代以降の西洋哲学の産物です。その基本には、「他行為可能性」といって、人間が「他の行為も選択できた」状態であることを前提に罪や責任を問う姿勢があります。しかし、この姿勢は必ずしも自明なものではない。例えば古代インドの発想を借りれば、「事故が起きたのは、前世からのカルマ（業）によるためだ」と説明することも可能です。または、古代ローマやギリシア時代にもし奴隷が罪を起こした場合は、その主人が管理不行き届きの責任を背負うことがあります。私たちの研究では、このように他行為可能性を前提としない過去の社会制度を参照としながら、他行為可能性や近代的な「責任」や「主体」から形成される社会が果たして本当に幸せな制度なのかも同時に議論しています。

どこからが機械か？人間か？

ー 人工知能側に責任が問われる場合、どのレベルで機械が「自律的である」と判断できるのでしょうか。

松浦： それは、人工知能が人々の日常生活に浸透したときにこそ、より重要になってくる問題だと思います。しかし、機械の自律性をレベル分けしたり、「自律性」の概念を定義するよりも、機械に自律性や意志があるように「見える」ことから議論を発した方が、人工知能をひとつの「文化」として社会になじませることができるのではないかと考えています。それゆえ、ご質問には、それぞれの文化が「自律性」をどのように捉えるかによって判断が分かれるだろう、と差しあたってお答えせざるを得ません。

葭田： サイエンスの視点から見ても、「アニメシー」など、どういう時に人間は機械やCGに知性や生命性を感じるかという研究があり、そちらの方が定義は簡単ですし研究もしやすいですね。しかし、それはチューリングテストにも通じる話で、その対象は人工知能というより「人工無能」かもしれない。つまり、たとえ生命がないかもしれないものにも人は生命性を感じられるともいえてしまうため、現在議論されているような徐々に高度化している

人工知能に対しても同等の議論をして良いかはよく考えなければならないですね。

稲谷： 法学の主流なやり方においては、人間が備えるべき責任能力の有無という観点から自律性を捉えようとします。しかし、先ほど申し上げたように、そもそも、法的責任を負いうる「人間とは何か？」という問いの答え自体が揺れているわけです。従来の法学は、近代哲学に依拠して、「人間とはかくあるべき」という地点からスタートし、そこから外れたものに対し制裁を加えてきましたが、そのやり方自体を再考する時期に来ているかもしれません。



松浦： 「自律的な人間」をまじめに追求していったら、外部から一切の影響を受けない人間、ということになってしまいます。そうだとすると私たち人間の中に果たしてそんな人が存在するのか疑問ですよ。強いと挙げるならカントくらいじゃないでしょうか（笑）。

葭田： カントさんの人間性までは個人的に存じ上げませんが（笑）、もし完全に外界から自立した人間や人工知能ができたら、きっとそれは社会通念も完璧に無視して行動しそうですから、そうすると、通常の社会で果たして無事に生きていけるのかなあという気はしますね。

技術的な「安全」と心理的な「安心」のズレ

ー 最近では、自動運転車による事故はメーカーと運転者のどちらに責任があるかが社会的な議論になっています。皆さんはどうお考えですか？

稲谷： 自動運転は完全自動化に至るまで0～5の技術レベルが定義されていて、それぞれのレベルによって条件が全く異なるので一概には言えませんよね。

葭田： 私は自動運転に関しては、レベル3程度の半自動運転（高速道路など特定の場所で機械が運転し、人間は運転席に座って緊急時のみ対応する）の設計思想があまり語られないことが気が

なっています。機械と人間が中途半端に協調動作をしている状態は、それぞれが単独で動作している時と比べると事故が起こりやすい可能性があってもおかしくなさそうな気はします。これは航空機のオートパイロット実験など既に他の自動化技術で経験的に知られていることかもしれない、完全自動運転の方が事故は起きにくく、中途半端に人間の操作が介在すると事故が起きやすいと主張する研究者は一定数いるそうです。

稲谷： レベル3については、最終的には人間が意志に基づいてコントロールした方がいいと無意識に想定されているからこそ出てきている話だと考えています。近代法からすると、ドライバーは人間だから危険を見抜けるだろうし、また人間である以上危険を察知し、コントロールする責任があるという発想はある意味素直です。そこから引き起こされる問題は甚大かもしれませんが。

葭田： 仮にどんな問題も完璧に解ける人工知能が電車や航空機を運転していたとして、それに人間は乗りたいと思うか、という視点もありますね。人間の普通の感情として、有事の際は機械よりも責任がとれる人間が操縦しているほうが安心できるのだと思います。

稲谷： 心理的な安心感と、客観的な安全性とのズレですね。そこが議論を錯綜させている一番の問題なのかもしれません。

葭田： 特に半自動車運転の場合は、人がどういう認知特性を持っているのか、運転時にどういう状態であれば、安全で快適かつ自分の責任感を維持できるかを踏まえた上で、そういう人間を内包して共に走る機械システムの姿とはどうあるべきかを踏まえた設計思想があっても良いように思います。

松浦： 機械と人間が協働するという点で言うと、機械システムが人間の行動や能力を拡張・またはサポートするよう設計するアプローチもありますよね。これからは私たちの能力をエンハンスする技術にも期待を寄せたいです。自動運転技術も、完全な自動化を目指すよりも、人間の運転能力を補佐する力を向上させるという方向で進んでもよいのではないのでしょうか。

問題を追及するのではなく、理想の社会ビジョンをつくる

稲谷： いずれにしても、今後は人工知能によって危険が引き起こされたからといって、利用者や開発者の責任ばかりを追求していくのはあまり生産

的でないと思います。今後は機械や人間の本質を問うよりも、まずどんな社会像を目指しているかを皆が考え、その目的に見合った法的責任の分配とはいかにあるべきかを議論すべきタイミングにきているのではないのでしょうか。自動運転であれば、どのような形で、どれくらい自動運転車が社会に浸透しているのがいいかを、皆が具体的にイメージすることから始めた方がいいと思います。そうした観点からすると、あるべき社会像を本質化して、事前に規制をかけていこうという従来のやり方は少しゆるめたほうがいいと感じます。社会はどうあるべきかの議論ではなく、どうありたいのかについての議論を、少しずつ生かしていくアプローチの方が望ましいと私は思っています。

葭田： プロトタイプを将来のユーザーに渡してみて、フィードバックを得ながら考えていく方法もありますよね。



稲谷： そうですね。実際に流通させた後も、それが何か事故や問題を起こした場合には、ユーザーや企業、専門技術者、さらに法曹関係の人たちが入った状態できちんと議論をして、「ここまでコントロールできるものならこうしよう」あるいは「ここまでコントロールできないなら、こうしよう」という具合に、みんなが目指す方向性を、その都度議論しながら、少しずつ理想の状態に近づけていく制度を用意したいと考えています。

人々の声をいかに取り入れるか

ー 人工知能が社会に受け入れられるには、理屈や制度では割り切れない課題も生じてくるのではないのでしょうか。

稲谷： 誰が主体となって議論するかが重要ですよ。例えば、今ならSNSを通して一般の人々から広く意見を取り込むこともできるでしょうし、その声を人工知能がデータとして解析することもあり得ます。専門家の議論に偏らず、一般の人たちの意見にしっかり耳を傾けることは重要です。一方で、ユーザーの意見に重きを置きすぎると不安の

声も高まり、「とりあえず危険は排除しよう」と余計に規制が厳しくなりかねない部分もあるので、それにも注意が必要だとは思いますが。

松浦： 人々の意見を社会制度に反映させることは必要な手続きだとは思いますが、「誰の意見を集約するか」と、「集約した意見がこれまでの法体系や文化と調和するか」という2つの問題が出てくると危惧しています。実際に人工知能の学習過程において、どんなインプットを与えるかによってその後は全く異なる反応を示すようになります。暴力的な人たちの中で学習をしたら、人工知能は暴力的な出力をするでしょうし、特定の宗教への関心が高い人たちを集めれば、その宗教における思想を反映させた出力をするようになるでしょう。同様に、一般からの意見の集約が重要なプロセスのひとつであることは否定しませんが、それをどうまとめ上げるかという手法も重要になると思います。そうでないと、地域や人の嗜好ごとに異なるばらばらなルールができるか、過度に理想化された人間像に沿った法律になってしまうと思いますね。

ー 最後に、皆さんのこれからの共同研究の可能性についてお聞かせください。

葭田： まずは異分野の人の話を丁寧に聞く姿勢や機会をどのように作るかが第一と思っています。私のいる東工大の、特に機械工学の学生や研究者たちは人文的な思考に接する機会が多くはないので、近々、松浦さんと稲谷さんを東工大にお呼びして、法学や哲学など他の分野と対話していくきっかけを作りたいですね。それに、この種の議論をもっと多くの人たちを巻き込んで続けていきたいですね。

松浦： 葭田先生と同様、私たち人文系研究者もテクノロジーの先端が今どういう状態にあるかをより深く知り、専門の先生とコミュニケーションできる機会をもっと増やしていくべきだと思っています。この機会を通じて、現代の専門分野の垣根を越えた議論をしていきたいですね。さらに今後は、こうした議論を教育の現場に展開する具体的な手段も考えたいと思っています。

稲谷： 教育でいうと、最近高校生向けの授業を行ったのですが、デジタルネイティブの彼らは、情報技術や人工知能の特性に対する理解もすごく早いし、それらがもたらしうる変化についてもとてもオープンに受け入れているんですよね。その感覚はとても重要だし、彼らによって人工知能がどん

どん社会で柔軟に使われていくことで見てくる可能性もあるでしょう。また、法学はどのような領域とも交わる分野なので、どこにでも流通できるという側面があります。他の専門領域の方々の知見を活かして、今後もより具体的な法提案につなげていければと思っています。



松浦： これまでの近代法が築いてきた枠組みを超えなければならない時代において、最善のものを提示はできなくとも、採用しうる選択肢くらいな哲学の人間が提示できるはずですし、その提示を前提として社会がどう変容するかという考察もできると思います。今日お話してみて、稲谷先生のような法学の方とはかなり通じるところがあると思いました。ようやく、私たちの多様な価値観や生活になじんだ法律や社会制度がつくれる時代が来たのではないかと。

稲谷： このHITEの取り組みを起点として、日本ならではの法のあり方を提示できるのではないかと思います。特に「主体」と「客体」の線引きを明確に維持しようとする伝統のある西洋文化圏では、今日のような「主体」と「客体」の曖昧さを前提とした議論を積極的に生かすのは、やや難しいように思います。そこにこだわる必要のない日本だからこそ、西洋近代哲学の枠に止まらない、新たな可能性を示すことが出来る部分はあると思うのです。このつながりを、世界を巻き込むような面白い取り組みに挑戦していくチャンスにしたいと思っています。

松浦 和也
HITE プロジェクト「自律機械と市民をつなぐ責任概念の策定」代表。秀明大学学校教師学部 専任講師。専門はギリシア哲学。
葭田 貴子
HITE プロジェクト「人間とシステムが心理的に『なじんだ』状態での主体の帰属の研究」代表。東京工業大学 工学院准教授。専門は応用脳科学。
稲谷 龍彦
HITE プロジェクト「自律性の検討に基づくなじみ社会における人工知能の法的電子人格」参加メンバー。京都大学大学院法学研究科 准教授。専門は刑事法。



想定外の未来に備える思考法

鷺田 祐一

**HITEプロジェクトにて「未来シナリオ」を用いた未来洞察の研究を進める鷺田祐一教授。
ものすごい速度で発展する情報技術を前に、新たな未来予測が求められる現在、
想定外な未来に備えるための、未来シナリオの可能性と意義を伺いました。**

ー 鷺田先生の研究では、AIと人間の共存などといった予測困難な未来において「未来シナリオ」を用いた思考法が有効だとされていますね。

日本の国家そして企業における、直近30年間の意思決定を振り返ると、その失敗の背景に共通項が見えてきます。たとえば先の30年間で日本が直面した想定外の事例として「パーソナル情報通信デバイス」と「第五世代コンピュータ」

の開発があります。1991年、当時の経済企画庁総合計画局の「2010年技術予測研究会」では、2010年段階での国内パーソナル情報通信デバイスの市場規模は約5000億円と予測されていました。しかし実際には、2005年の段階で市場規模は約2兆円を超え、政府の予測を5倍以上も上回る成長を遂げました。これはいわば上振れの想定外です。

一方で、第五世代コンピュータの開発は下振れの想定外。1982年、当時の通商産業省（現・経

済産業省）は、「世界に先駆けてAIを搭載したコンピュータの開発を目指す」として、第五世代コンピュータの開発を国家プロジェクトとして始動しました。しかし10年後、ほとんど成果を挙げることなくこのプロジェクトは終了します。非現実的な目標に投じられた資本は総額570億円にのぼりました。

当時の目論みはなぜはずれたのでしょうか？それは日本の国家や企業が高度経済成長期から一貫して、技術的な視点のみで未来を予測してき

たから、と言えるでしょう。

対話型のワークショップを軸とする未来シナリオづくりは、そのプロセスにおいて、技術的な視点はもちろん、人間における文化的背景、人口構造などの社会側面や自然環境の変化など幅広い視点を取り入れています。あえて政治や文化などの不確実性の高い事象についても積極的に議論することで、技術発展中心の線形な予測からではない、オルタナティブな意思決定の材料を提供できると考えています。

非線形で描く、 広い視野の未来像

ー 未来シナリオのワークショップは
どのようなプロセスで行われるのでしょうか？

RISTEX HITEプロジェクトで実施した「第1回未来洞察ワークショップ」では、これからのAIやIoTの普及に関する未来シナリオを作成しました。参加者は主にマーケティング領域における有識者で、2日間かけて行われました。初日は、10～20年先に起こり得る社会変化をシナリオ化する「社会変化仮説」を「ホライゾン・スキャンニング」という手法を使って複数作成します。まず未来の変化の予兆となるようなデータベース「スキャンングマテリアル」を用意し、それらをもとに参加者と議論を深め、これからどんな社会になるかをイメージしていきました（次ページに構想された社会変化仮説の「未来年表」を掲載）。2日目には、AIやIoTの普及が今後普及したときに起こりうる問題を「未来イシュー」と称して考えていきます。このときは、自動運転車と従来の自動車の混在による軋轢と二極化などの例が挙げられました。次に2日目で生まれた社会変化仮説と未来イシューを組み合わせるさらに議論を深め、最終的に「未来シナリオ」を作成していきます。

ー これまでの未来シナリオ研究の長い経験の中で、興味深かった未来シナリオの事例を教えて下さい。

2002年に、KDDIと「2008年における秋葉原の未来像」を描く未来洞察ワークショップを実施した際、「メガネ型携帯電話」が発想されたのはいま振り返ると興味深いですね。このときから既に位置情報を活用したウェアラブルデバイス

が想定され、またそれに伴うプライバシーの問題が議論されました。それから十数年の後によく似たシステムの「Google Glass」が登場しましたが、やはり開発側も個人情報の問題が解決できず、民生品を諦めてBtoBに特化しました。未来を予測した示唆的な事例のひとつでした。

2025年問題と、 モザイク型社会の到来

ー AIなどの急速な情報技術の発展は、
今後ますます予測が困難な未来に直面
していきます。未来洞察には何ができると
お考えでしょうか？

何度もワークショップをやっていくと、今後の社会変化仮説に共通のパターンがあることがわかってきました。それは2025年を境に、社会が楽観的なものから悲観的なものに転じるということです。構成メンバーやテーマを変えても同様の現象が生じるので、多くの人が「2025年に想定外の事象が発生する」と考えているようです。この現象を私たちは「2025年問題」と呼んでいます。

そこで私たちはこの「人と情報のエコシステムプロジェクト」にて未来洞察ワークショップを開催し、「AIやIoTにおいても同様の2025年問題が発生するのか」というテーマで検証を試みたところ、新たな仮説が生まれてきました。それは、社会が変化し続けていくとき、人々の中で必ずどこかに破綻が起きるというもの。私たちはその現象を「モザイク化する社会」と名付けました。「モザイク化」とはAI技術の実装が社会の中で進展するゾーンと旧態依然としたまま残るゾーンがモザイク模様のように入り組んでくるという意味です。現在よく言われる、AIが人間の仕事を奪うか、支配するかといった単純な未来予測ではなく、なってくるだろうと。企業や開発者の間では技術が線形に進展するような理想像が描かれるばかりで、こうしたモザイク型の普及などは想定されていない。そのギャップを埋めるのが未来シナリオだと思っています。

ー 未来シナリオは、
日本の未来にどのような貢献をするのでしょうか？

未来シナリオは、適切に利用されれば社会に

おける予測困難な状況に対する「備え」をもたらしてくれるものです。特に情報技術に携わる企業にとって、AIやIoTなど情報技術がどのような形で社会に浸透するかを考える上で有効な手法だと思います。



従来型の発想では、事業計画を進める際に「調査を行い、現在の技術動向からトレンドを洗い出そう」と考えがちです。しかし情報技術の多くはコンシューマーエレクトロニクス、つまり消費者にとって非常に身近な技術です。

現在のスマートフォンに代表されるような情報技術には「ネットワーク外部性」（サービスの利用者の増加が、サービスそのものの価値を向上させる効果）が働くため、あっという間に他の追随を許さない製品になり得ます。つまり先行者利益が非常に大きいため、参入する事業者が「先行逃げ切り型」になるのです。またその新しいサービスの多くはアメリカからの流入だったりもしますね。

こうした特性を持つ情報技術においては、あらゆる条件で起こり得る未来を想定しておかなければなりません。そうした状況下において、未来シナリオの思考法は有効だと考えています。

鷺田 祐一

HITE 採択プロジェクト「未来洞察手法を用いた情報社会技術問題のシナリオ化」代表。一橋大学大学院商学研究科教授。研究分野はマーケティング、イノベーション研究、未来洞察など。

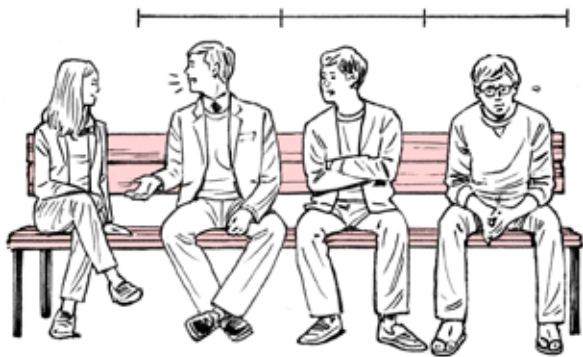
第一回未来洞察ワークショップ

AIはどう社会に普及するか？

2040年までの未来ロードマップ

2018

若者は3種類に分類できる



- 若者はコミュニケーション・エリート、中間層、コミュニケーション困難層の3つに分かれ、その間の断絶が顕著になる。
- 若者の年齢による横のつながりが弱くなり、コミュニケーションにおける個人主義化が進行する。
- シェアリングエコノミーが充実し、ライフスタイルは、すべてネットワーク上で管理する「クラウド化」、もしくは貯蓄・所有を行わない「フロー化」が進む。
- 学校が唯一のコミュニティハブとして機能するようになる。
- 断絶を感じる人々向けのビジネスが、「知識」から「癒し」へシフトする。

2020

裕福な高齢者が
“あしながおじさん化”
する



- 極端な核家族化の結果、お金の余った裕福な高齢者が血縁に関係なく若者を経済的に支援する＝あしながおじさんとして社会に影響を与える。
- 寿命が伸長し、若く生きる時間よりも、老いて生きる時間が長くなり、多くの人が「老後のために生きる」人生を送るようになる。
- スマートデバイスのインターフェイスがより簡易なものになり、AIのユニバーサルなサービス化が進む。
- 商品やサービスを自ら探すのではなく、AIによって自動的にマッチングが行われる「サービスマッチング社会」が到来する。

2020

国の在り方が
変わる

- 年金制度の破綻や医療費の自己負担の増加に端を発して、これまで行政が国民に提供してきた福祉が限定的になる。結果として、国家が市民に与える影響力が弱まっていく。
- 国家を代替するような、多様なコミュニティが形成。既存の国家から独立する機運が高まる。
- 既存の国家システムとの間に摩擦が起きる。また、その摩擦、情報のアービトラジーを利用したビジネスが創造される。



2030

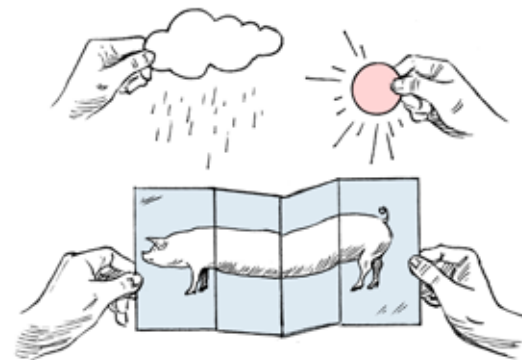


AIが“モザイク社会”を
形成する

- 産業社会生活資本にAIが導入され、AIに決断「任せる」人と、自分自身の決断に「こだわる」人が社会の中で混在し、“モザイク化”していく。
- 現実世界のビッグデータを解析し、意味あるデータを抽出して活用する技術が飛躍的に向上する。サイバー世界と、フィジカルな現実世界が緊密に連携する、社会の「サイバーフィジカル」化が進行する。

2035

人間が自然を自由に操作できる
ようになる



- 農産物や畜産物などの遺伝子組み換え技術が進み、またその育成環境も人工的に構築されるなど、自然そのものが人為的につくられるものになる。
- ある種の気象操作などが可能になるなど、自然のコントロールによって新たに持続可能なエネルギー資源を手に入れる。
- 人為的に食糧をつくりだす技術が確立し、普及。食品革命が起きる。

2040

バーチャルおよび
リアルの世界で
民族大移動が
起きる



- 地球の気候変動によってニューノース時代（高緯度地域の繁栄）が到来。リアルの世界で民族単位、または個人間での移住が活性化する。
- 同時にクラウド上のバーチャル国家も繁栄する。それらが持続的社會変革を牽引する。
- 情報技術の発展によるテロや戦争、またはバイオテクノロジーの過剰な進展による未知の事故や病気といった想定外の負の影響もたらされる。世界規模のリスクが増大し、それらを活用していく社会が誕生する。

AI 技術は
モザイク型に
普及する

未来洞察の手法を研究・実践する鷲田祐一教授らは、HITEプロジェクトにおいて、有識者を招いてAIやIoTなどの情報技術が普及した未来を予測する「第一回未来洞察ワークショップ」を開催した。その結果、2025～30年頃にAIやIoTの「モザイク型普及」が起きるという未来シナリオが描き出された。その背景には、社会格差が進み、国家による福祉サービスが弱まるという予測がある。結果として、「AIが人間の仕事を奪う」、「社会が一斉にIoT化する」といった、昨今よく話題となる画一的な変化が起こるのではなく、テクノロジーの浸透は社会の中でかなりの“ばらつき”を持って進むというシナリオが生まれた。このばらつき現象を鷲田教授らは「社会のモザイク化」と名付けたが、その問題を見越した上で、今後のテクノロジー発展とどう向き合うかがこれからの大きな課題となるだろう。あらゆる未来を想像し、柔軟に備えをしていく思考のトレーニングが「未来洞察」なのだ。

※同ワークショップの結果内容は『マーケティングジャーナル Vol.37 No.1(2017)』にも掲載されている。

AIと社会の関係を「^{べきそく}冪則」で読み解く

田中 久美子

人間のコミュニケーションや自然言語における普遍性を数理的に捉える研究に取り組む田中久美子教授は、「AIがいかに社会へなじんでいるか」を評価する軸として、「^{べきそく}冪則」という統計モデルの観点を採用しています。その研究の意図と、目指すべき社会ビジョンを尋ねました。

田中久美子
東京大学先端科学技術研究センター 教授。専門
は計算言語学、コミュニケーションの複雑系科学、
自然言語処理など。HITE 採択プロジェクト「冪
則からみる実社会の共進化研究－ AI は非平衡な
複雑系を擬態しうるかー」* 代表。

この世界のあらゆる場所^{べきそく}で生じる「冪則」

－「冪則」とは
どんなモデルなのでしょうか。

冪則とは、生物や地震、経済現象、自然言語、または都市人口の分布など、さまざまな統計やシステム上で経験的に観測される物理法則です。基本的には(定数項などを除き)、2つの変量の間に両対数軸上で 比例関係が成り立つモデルのことを指します。

不思議なことに、これは様々な対象で観測される現象なのです。例えば、地震の統計では横軸にエネルギー、縦軸に頻度を並べるときれいな冪になりますし、都市の大きさに対する人口や、論文がリファレンスされる数と頻度をプロットすると、やはり冪になるのです。長者番付で、順位に対する年収をプロットして両対数を取ってみたら、線形に真っすぐな線が現れたりもしますね。

－ ある「冪則」が成り立つときは、
どんな場合でも例外なく成り立つのですか？

それが驚くべきところで、例外なく成り立つ場合が多いです。私の専門は元々自然言語なのですが、文章の中に含まれる単語の統計を取ると、日本語の場合は「の」が最も多いことが分かります。その次が「に」や「て」です。単語の頻度を順序に対してプロットしていくと、やはり冪則グラフになります。それは古今東西、老若男女、どんな言語の文章をサンプルにしても例外がなく、3歳の子どもが話す言葉にも当てはまります。そうした結果を意識して言葉を話す人は誰ひとりいないと思いますが、不思議なことにそうになってしまうんです。そして、それ

がなぜかについては、まだ解明されていません。

－ この「冪則」とAIはどうつながるのでしょうか。

機械はシンプルに設計されているので、冪則が成り立つかは自明ではありません。例えば10年前まで、機械が生成する自然言語には、人間の言語と同様の冪則はまったく成り立ちませんでした。Twitterのbotは自動的に誰かの言葉を引用していますが、あれでも冪則になるかは疑わしいですね。ただ最近では、ディープラーニングを用いると「傾き－1の冪則」(*fig1参照)が成り立つようになりましたが、全体的にはまだまだです。

それが今回のプロジェクトの発端であり、「AIの導く結果がどれだけ自然な状態に近い」すなわち「どれだけ社会になじんでいるか」を見定めるための尺度として、冪則が成り立つかという問いを取り入れようと提案したのです。

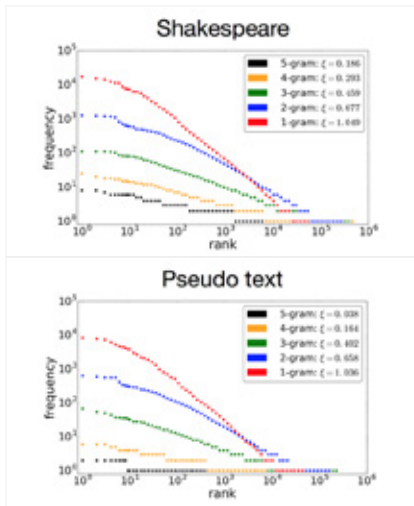


fig1: [上] シェイクスピア全文のべき則。[下] それを学習したディープラーニングが生成した文書のべき則

－ 冪則の成立については、
何が指標になるのでしょうか。

冪は基本的に、自然界で生じるコッホ曲線などの「フラクタル構造」に類する統計的自己相似性を表します。自己相似性とは、ある仕組みの中に小さな同じ仕組みがあって、さらにその中に同じ仕組みがあるという入れ子構造になっていることですが、冪もまた、どれだけ拡大・縮小してもやはり同じ線形の傾きが現れます。これは多くの自然や社会的対象に見られる性質ですが、機械は人工的なので、基本的にはそうした構造がない。ですから、機械が発生する系の中に、自己相似性があるかを測る指標として冪則を見ていきます。

冪則は、AIを評価する新たな指標になる

－ その尺度をAIに取り入れると、
どんなメリットがあるのでしょうか。

冪則が成り立っていないAIシステムを検出することで、例えば投資分野で危ない暴走行動に走るAIを排除する仕組みなどが作れないかと考えています。株式投資においては、その時々で儲かることのみを判断していくと、おそらくAIは同じ行動を取り続けるでしょう。そうするとあつという間に冪ではなくなり、リスクが高まる。その兆候を冪則で判断できるのではないかと考えています。

また冪の考え方は、AIが導く判断基準に多様性を取り入れることにもつながります。人間の社会には「ロングテール」という考え方がありますね。一昔前まで、映画や音楽などのコンテンツビジネスはヒットさえ出せばいいと考えられてき



ましたが、ヒット作品のチャートだけだと冪にはならず、必ず売上に限界が生じることがわかっています。例えば1万曲を有するデジタルミュージックボックスをブロードバンドでインターネットに接続して約3カ月間稼働させたとき、なんと98%の曲がお客からセレクトされるのだそうです。人の好みは多様なんですね。その裏付けはECサイトなどでますます顕著になってきています。ECショップが人気商品だけ在庫を増やせばいいわけではなく、売上結果を見ると実は少人数のユーザーが利用する多様な商品の在庫が実際の経営を支えていることがわかります。それがグラフ上ではなだらかに続く長いしば（ロングテール）のような形になるのです。

－ 今回のプロジェクトでは、具体的にどんな
取り組みをされていくのでしょうか。

大きくは技術研究と社会実装という2つの取り組みがあります。前者では、AIの導く統計において冪があるかどうかをまず分析します。大抵の場合は成り立っていませんし、仮にできていても

大きな問題があります。そこで、どんな数理モデルを与えると、AIにも冪則が生まれるかを考えていきます。後者の社会実装では、AIによる投資行動の抱える問題を明らかにし、社会制度の提案につなげることに向けて動き始めています。



－ AIの投資行動には、
いまどんな問題があるのでしょうか。

近年、AIが株式市場に参入したことを受けて、2017年5月に金融商品取引法が改正されました。しかし、その内容の主な部分は、常にログを取っておき、問題発生時にはその記録を金融庁に提出するというものなのです。それって、事後

対策に過ぎないんですね。例えばAIが原因でブラックマンデーのような株の大暴落が再び起きたとして、いくら後から原因を解明できてもその後の社会に与える多大なる影響を考えると手遅れです。ですから、何かが起きる前、少なくとも起きている途中で、それを食い止める手立てがないとまずい。そこで、冪則をひとつの頼りに挙動を評価できないかと考えています。やたらと利益ばかりを得ようとするAIをどうやって事前に食い止められるか、その時の評価方法をどう設計するかを議論しています。

このプロジェクトをきっかけに、経済活動やその他さまざまな現場にいる方に対しても、冪という側面からAIを評価して、社会の健全さを維持するための切り口を提供していきたいですね。

*JST 戦略的創造研究推進事業（さきがけ）にて、2014年度から展開してきた基礎研究が同プロジェクトにつながった。

PROJECTS OF HUMAN-INFORMATION TECHNOLOGY ECOSYSTEM

01

人間とシステムが心理的に「なじんだ」 状態での主体の帰属の研究

ヒトと協調して自律的に動作可能な機器やシステムが、事件や事故など社会的に思わしくない行為を引き起こした際、それはユーザであるヒトが主体的に引き起こした行為でありヒトが責任を追うべきという考え方と、機械やシステム側が主体的に起こした行為であるためそれらの製造者側が責任を負うという考え方のどちらを採用すべきか、その科学的な考え方や解釈を心理学や脳科学の立場から提案する。特にユーザであるヒトからみて、その事件や事故が機械やシステムではなく自分自身が引き起こした行為と錯覚され、不必要に責任を負う状況の存在を指摘し、そのような錯覚や思い込みの背後にある脳科学的仕組みの解明や、そのような錯覚を逆手に取った機器の開発とデモを実施しながら、我々人間自身ですらどこまで自律的・主体的に行動する存在といえるか考察する。



稲田貴子
東京工業大学工学院・准教授
- 心理学 -

02

自律機械と市民をつなぐ 責任概念の策定

AIを搭載した自律機械の社会実装が現実味を帯びつつある中、社会的不安も同時に発生している。その背景には自律機械が起こした事故や、生み出した被害への責任の帰属先が不明瞭であることが理由にある。この不安を軽減する手段として、本プロジェクトでは、自律機械が実装された社会の中でも説得力を持つ責任概念を提案する。自律機械が社会になじむには、その社会的位置づけが非専門家たる市民にも納得できるように、歴史的・文化的背景からも説明されなければならない。そのためには、現在の技術的・社会的状況を踏まえつつも、「自律機械が人間と対等と見なされるためには何が必要か」という問い（それは「人間とは何か」の裏返しでもある）を人類が積み上げてきた人文的知見に投影する、哲学的考察が必要である。この考察に基づき、自律機械と市民をつなぐ新たな責任概念の策定を目指す。



松浦和也
秀明大学学校教師学部・専任講師
- 哲学 -

03

自律性の検討に基づくなじみ社会 における人工知能の法的電子人格

近年のAI技術は人工システムやロボットのある種の自律性を可能にし、丁度、親離れた子供のように、設計者が予測できない行動を表出する可能性がある。このような状況に対し、現在の法制度では設計者が利用者が過度の法的責任を負わされる恐れがあるため、健全な科学技術の進展を阻害する可能性がある。本プロジェクトでは、人工システムの自律性を目的の有無やその書き換え可能性に準じて、三段階を想定し、従来の法人格論の分析を通じてこの三段階に応じた法的取扱モデルを考案する。さらに、既存の責任理論の問題点を指摘し、人工システムに対する新たな制度の提案を行う。また、アンドロイドを用いた模擬裁判を通じ、一般社会になじんだ法整備案を提案し、自律性の概念の深化と未来社会に通用する人工システムとその環境を提示する。



浅田稔
大阪大学大学院工学研究科・教授
- ロボット学 -

04

情報技術・分子ロボティクスを対象とした 議題共創のためのリアルタイム・テクノロジー アセスメントの構築

本プロジェクトでは、先端情報技術の倫理的・法的・社会的影響（ELSI）について、メディア分析と予測評価手法（ホライズン・スキャンング）による議題抽出を行い、さらに先端情報技術の専門家をはじめとした多様なステークホルダーが参加する「議題共創プラットフォーム（NutShell）」の開発を通じて、当該領域の社会的議論を迅速に焦点化する「リアルタイム・テクノロジーアセスメント（RTTA）」システムの構築を行う。このRTTAシステムを用いた議題構築について、分子ロボティクス分野ならびにAI分野の事例を通じた実践から、その課題抽出を行い、ELSIに関するより良い議題構築プロセスの実現と知見の現場の研究者へのフィードバックの在り方を提案する。



標葉隆馬
成城大学文芸学部マスコミュニケーション学科・専任講師・科学技術社会論・

05

分子ロボット ELSI 研究とリアルタイム 技術アセスメント研究の共創

分子ロボットの倫理的・法的・社会的課題（ELSI）の研究とインターネットを活用して技術・社会双方の幅広い知見・意見を集めるリアルタイム技術アセスメント研究をスパイラル的に推進することで分子ロボット技術と人間のなじみのとれている社会の実現を目指す。本プロジェクトでは(04) 標葉隆馬グループと共創し、リアルタイム技術アセスメント技術を活用した社会からの幅広い意見の集約、その意見を反映した分子ロボットガイドライン案の策定、策定案に関する研究会・シンポジウムによる議論、のプロセスを繰り返すことで分子ロボット ELSI 研究を進める。また、分子ロボット若手研究者および学生の理解を促進するために、分子ロボット国際学生コンテストに参画する。



小長谷 明彦
東京工業大学情報理工学院・教授
- 知能情報学 -

06

幕則^{（べきそく）}からみる実社会の共進化研究 - AI は非平衡な複雑系を擬態しうるか -

本来的にブラックボックスである今日のAIの社会への「なじみ度合」について、幕則（べきそく）の観点から評価する方法を研究する。評価方法を社会実装する一事例として、国の将来を担う投資分野における高度IT利用を前提とする適切な制度環境について、コミュニティを構築して議論し、法制度の方向性についてシナリオ分析を行う。テーマを技術研究と社会実装の2に分け、それぞれの視点から共進化プラットフォーム事例を形成する。技術研究ではAIが生成する擬似データと言語ならびに経済の実データとの幕則（べきそく）の観点からの差異を定性的に調査し、両者の乖離を明確にする。その知見に基づき、社会実装ではAIを投資に利用する場合の問題を社会的観点から議論し、限界をふまえてAIを適切に活用するための社会制度設計に向けて提言を行う。



田中（石井）久美子
東京大学先端科学技術研究センター・教授
- コミュニケーション科学 -

PROJECTS OF
HUMAN-INFORMATION
TECHNOLOGY
ECOSYSTEM

07

情報アクセスリテラシー向上のための
不利益の視点からの方法論に関する調査

システムを利用して情報を取捨選択し、批判的に意思決定をする能力を「情報アクセスリテラシー」と呼び、その向上を目的とする。効率や利便性を追求した情報アクセス技術は、人間の認知能力を低下させる側面も備えている。そこで、人が能動的に手間をかけ、頭を使うことでしか得られない効用である「不利益」に注目する。これを活用したリテラシーの維持・向上をめざして、意思決定支援システムの利用機会の増加が見込まれる医療従事者を対象とした調査を行う。具体的には、ヒアリング・アンケート調査による医療現場における情報アクセリステラシーに関する問題・ニーズの深掘り、リテラシー診断ツールの開発、医療従事者とのワークショップを通じた不利益受容可能性調査を実施する。

 川上浩司
京都大学デザイン学ユニット・特定教授
- システム工学 -

08

人と AI システムの協働タスクモデルの
構築に向けた調査

労働経済学、サービスマネジメント、知能情報学の3つの分野の知見を有機的に結び付けながら、先行導入・実験事例へのヒアリング調査やインターネットを通じたアンケート調査を実施。新たな情報技術や人の担うタスク（業務）の種類や量を定性的・定量的に把握するとともに、情報技術と人間・社会との協働・競合関係で留意すべき点を洗い出し、ビジネスや制度・政策における対応方策や望まれる方向性を提示するための方法の精緻化を図る。そうした分析を通じて、新たな情報技術の先行導入・実験事例をフィールドとした調査・分析と、国民全体あるいは消費者・労働者への量的調査・分析の2つを軸とした研究を実施するための準備を進める。

 山本 勲
慶應義塾大学商学部・教授
- 労働経済学 -

09

人工知能と労働の
代替・補完関係

AIの発達が我々の仕事を奪うという懸念が語られて久しい。たとえば、2015 年末出版のレポートには日本の雇用の49%が機械と代替可能であるとしている。これらの研究は従来の職業データベースの労働特性の軸に従って、労働とAI技術の代替補完関係をとらえており、AI技術の本質をとらえたうえで測定するという枠組みにはなっていない。本プロジェクトでは人工知能に代表される機械と労働の代替・補完関係を決定する根源的な原因を概念化し、それをサーベイの質問項目でとらえる方法を開発する。その際、カギになる概念は大規模な電子データの存在の有無と、因果関係の把握の困難性にあると考えている。この概念を科学者・エンジニアへのインタビューを通じて洗練したうえで大規模サーベイを実施し、既存のものを越えた新しい職業データベースを開発する。

 川口 大司
東京大学大学院経済学研究科・教授
- 労働経済学 -

ADVISOR

研究開発領域総括・アドバイザー
(2018 年 1 月現在)

総括：	國領 二郎	慶應義塾大学 総合政策学部 - 教授
総括補佐：	城山 英明	東京大学大学院法学政治学研究科 - 教授
アドバイザー：	加藤 和彦	筑波大学大学院システム情報工学研究科 コンピュータサイエンス専攻 - 教授
アドバイザー：	久米 功一	東洋大学経済学部 - 准教授
アドバイザー：	河野 康子	一般財団法人日本消費者協会 - 理事
アドバイザー：	砂田 薫	国際大学グローバル・コミュニケーション・センター - 主幹研究員
アドバイザー：	土居 範久	慶應義塾大学 - 名誉教授
アドバイザー：	西垣 通	東京経済大学コミュニケーション学部 - 教授
アドバイザー：	信原 幸弘	東京大学大学院総合文化研究科 - 教授
アドバイザー：	松原 仁	公立はこだて未来大学 - 副理事長
アドバイザー：	丸山 剛司	中央大学 理工学部 - 特任教授
アドバイザー：	村上 文洋	株式会社三菱総合研究所 社会 ICT イノベーション本部 ICT・メディア戦略グループ - 主席研究員
アドバイザー：	村上 祐子	東北大学大学院文学研究科 - 准教授

見守り技術の実装のための
現場変容ライブラリの構築

IoT や AI 技術の発展により見守り技術が高度化しており、保育園や介護施設など見守りが必要とされる現場での課題を解決することが期待されている。しかし、現場では期待とともに、プライバシー侵害のリスクや見守り技術の信頼性への不安があり、導入が進まない実態がある。また、見守り技術開発者にとってはニーズが分かりづらく、受容可能な状態も分からないため、開発に踏み切らないという状態である。本プロジェクトでは、現場の課題と見守りニーズを整理した上で、現場の見守り技術に対する受容範囲や受容可能な状態を明らかにする取り組みを行う。見守り技術を現場に実装するための方法論として整理し、見守り技術を必要としている現場への実装が促進されることを目指す。

 北村光司
産業技術総合研究所人工知能研究センター・主任研究員 - コンピューターサイエンス -

募集要項

研究開発の目標

「人と情報のエコシステム」は、科学技術振興機構 (JST) 社会技術研究開発センター (RISTEX) が推進する研究開発領域です。本研究開発領域は、ビッグデータを活用した人工知能、IoT、ロボットなどの情報技術を、人間を中心とした視点で捉えなおすこと、そして一般社会への理解を深めながら、技術や制度を協調的に設計していくことを目指します。情報技術と人間のなじみがとれている社会を目指すために、情報技術がもたらすメリットと負のリスクを特定し、技術や制度へ反映していく相互作用の形成を行います。

1. 情報技術がもたらしうる変化（正負両面）を把握・予見し、アジェンダ化することで、変化への対応方策を創出します。

2. 情報技術の進展や各種施策に対し、価値意識や倫理観、また現状の制度について検討し、望まれる方向性や要請の多様な選択肢を示します。

1、2 のような、問題の抽出、多様なステークホルダーによる規範や価値の検討、それに基づく提示や提言までをサイクルとみなし、その確立のための研究開発を行います。また、このような社会と技術の望ましい共進化を促す場や仕組みを共創的なプラットフォームとして構築することを目指し、その機能のために必要な技術や要素も研究開発の対象とします。

研究開発のアウトプット

領域全体として領域終了後には、技術と社会の共進化プラットフォームを構成する、以下の 1 ～ 5 の具体的なアウトプットが創出されることを想定しており、研究開発プロジェクトはこれらのアウトプットへの貢献が期待されます。

1. 情報技術の開発に社会的要請をフィードバックするための方法論

2. リテラシー向上のための方法論

3. 技術進歩に政策立案者やビジネスモデル設計者が対応して制度設計・マネジメントを行う仕組み

4. 技術と社会の対話に参加する人材のコミュニティ

5. 技術と社会の対話の共通基盤となる概念の構築

問合せ先

国立研究開発法人
科学技術振興機構 (JST) 社会技術研究開発センター (RISTEX)
「人と情報のエコシステム (HITE)」研究開発領域事務局
〒102-8666 東京都千代田区四番町 5-3 サイエンスプラザビル 4F
TEL: 03-5214-0133 / FAX: 03-5214-0140
E-MAIL: info-ecosystem@jst.go.jp

応募方法

「人と情報のエコシステム」領域サイト
<http://ristex.jst.go.jp/hite/> をご確認ください。

募集開始

平成 30 年 4 月半ば以降を予定。

Human-Information Technology Ecosystem Vol.02

Management:

Creative Director, Editor:

Art Director, Designer:

Writer (P.3-6,11-12):

Writer (P.7-10):

Photographer:

Illustrator:

Texts of cover visual*:

Published by

茅 明子, 芳賀 健一, 上原 ひとみ (JST)

塚田 有那

長谷川 弘佳 (H.Inc.)

高橋 ミレイ

森 旭彦

萩原 楽太郎

竹田 嘉文

浦川 通 (Qosmo)

科学技術振興機構 (JST)

社会技術研究開発センター (RISTEX)

* 本誌の表紙ビジュアルは、プログラマ浦川通氏の開発した言語解析プログラム「意識の辞書」を用いて、HITE に関するテキストデータの解析したデータからデザインしています。

