

NII Today

National Institute of Informatics News

103
Sep. 2024

[特集]

競争より協創

産官学連携による共同研究の現在

P2 日本メイトのLLMを構築する
黒橋 禎夫

P8 研究開発用LLMが果たす役割
鈴木 潤／横田 理央

P12 LLM構築における透明性・信頼性・安全性
関根 聡／宮尾 祐介

P16 人工知能法学とLLM
佐藤 健／西貝 吉晃

P20 自動運転車の「安全」の未来を協創する
上田 直樹／高尾 健司／吉岡 透／柳澤 名由太
／蓮尾 一郎／石川 冬樹

P24 工場をスマートにする無線通信
河村 憲一／金子 めぐみ

P28 サプライチェーンリスクを予測AIで可視化
佐藤 遼次／水野 貴之

P32 より自然で高強度な「音声匿名化」を求めて
高石 真美子／山田 正幸／山岸 順一

P35 アカデミアとインダストリーのストーリー
安里 彰／五島 正裕

Essay 日本の科学技術の明るい未来のために
片岡 洋



特集

競争 より 協創

競争、そして協創。

研究において、いつもさまざまな「競争」が繰り広げられている。そうして、最先端の科学はさらに研ぎすまされていくのかもしれない。

では、「協創」する利点はなにか。それは、研究者や組織、各々が持つ個性や専門性、技術、技能、アイデアを持ち寄り、響き合い、豊かなハーモニーとなり、新たな創生に至ることではないだろうか。

本号では、国立情報学研究所で進められている、そんな「協創」研究を特集する。

A portrait of Kurohashi Sadao, a middle-aged man with dark hair, wearing a dark suit, white shirt, and a dark tie with a small pattern. He is looking directly at the camera with a neutral expression. The background is a light gray with a faint geometric pattern of circles and lines.

黒橋 禎夫

KUROHASHI, Sadao

国立情報学研究所 所長

同 大規模言語モデル研究開発センター長

聞き手

吉川 和輝

YOSHIKAWA, Kazuki

日本経済新聞社 編集委員

(敬称略)

日本メイドの 大規模言語モデル LLMを構築する

文部科学省の「生成AIモデルの透明性・信頼性の確保に向けた研究開発拠点形成」事業を実施する拠点として、2024年4月、国立情報学研究所(NII)内に「大規模言語モデル研究開発センター」(LLMC)が設置された。大規模言語モデル(LLM)の透明性、信頼性の確保に向けた先端的な研究開発を行うことをミッションとする、センターの体制、今後について、同センター長の黒橋禎夫NII所長に聞く。

LLMCの新設

——新設した大規模言語モデル研究開発センター(LLMC)にどのような思いを込めていますか。

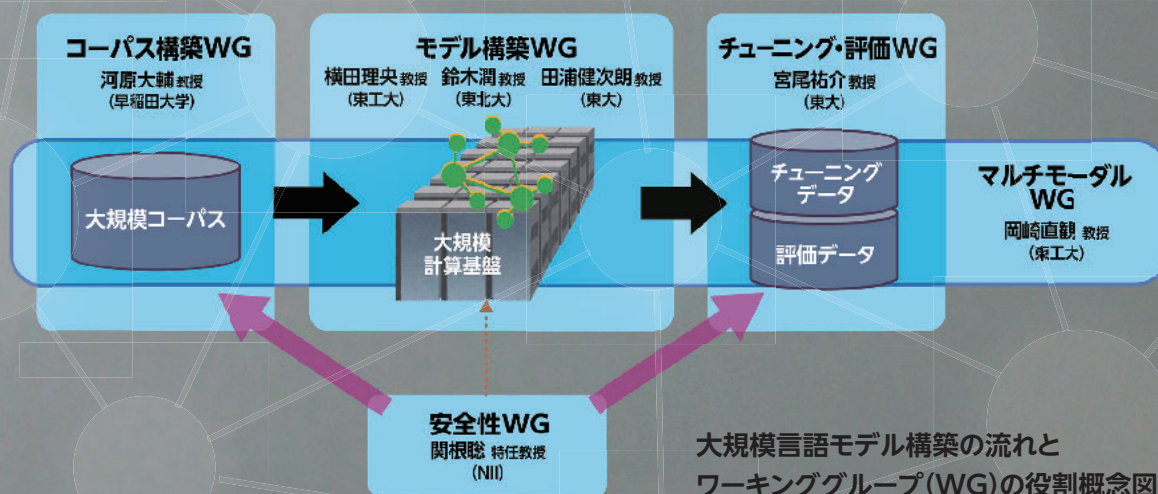
言語や画像などを生成する生成人工知能(AI)が社会にもたらすインパクトは極めて大きいものがあります。私は大規模言語モデル(LLM)こそが今の生成AIの本質であると考えています。LLMによって人間とAIの言葉によるコミュニケーションが

可能になったためです。LLMをベースにしたOpenAIのChatGPTは、人の言葉を理解してその意図を推測する能力が相当程度あります。人間にとって参考になることを言葉で示してくれるなど、生成AIとの対話を通じてコミュニケーションを深めていくことができるようになりました。

人間と同じように話をしたり、立ち振る舞いができたりするAIエージェントやロボットはSFの世界のもので、それが実際に登場するのはかな

り先になると思われていたのが、LLMを軸にした生成AIによってこうした世界が実現しつつあります。

この分野の技術進展のスピードにも目覚ましいものがあります。変化が急速過ぎると感じるほどです。生成AIの世界で何が起きているのか、どんな将来見通しが立てられそうなのか、またどんなリスクを想定しなければならないのか。これらについて正確な認識を持つことが不可欠です。LLMCはLLMの研究開発を



通じてこうした社会の要請に答えていきます。

—— ChatGPT が登場したのが2022年11月末でした。その数カ月後にNIIが研究者を集めて組織した「LLM勉強会」がLLMCの前身になったと聞いています。

ChatGPTの登場を受けて、まずはLLMの勉強をしないことには始まらないという話になりました。自然言語処理やコンピュータ科学の研究者に声をかけて、2023年5月にLLM勉強会を立ち上げました。当初30人くらいの規模で始まり、1年あまりたった今では参加者は1,600人を超える規模になりました(2024年7月現在)。「LLM-jp(エルエルエム-ジェイピー)」という名前で活動を続けています。

LLM-jpでは2023年10月に最初のLLMである「LLM-jp13B」を、2024年4月には改良版の「LLM-

jp13B v2」を公開しました。モデルの規模を表すパラメータ数は130億で、ChatGPT無料版で当初から使われていたLLMである「GPT-3.5」(推定1,750億パラメータ)の10分の1くらいです。

今年(2024年)4月にLLMCが発足し、LLM-jpの活動と連携しながら、パラメータ数1,720億、つまりGPT-3.5にほぼ匹敵する規模の「LLM-jp172B」の開発を始めました。学習に一部不具合がありましたが、その後は順調に進んでおり、2024年中に完成・公開の予定です。

その後は予算が認められれば、さらにその10倍、1兆を超えるパラメータのLLMの開発に取り組みたいと考えています。トップクラスのLLMであるOpenAIの「GPT-4」のパラメータ数は1兆を超えると推定されており、これに匹敵する規模のLLM開発への挑戦になります。

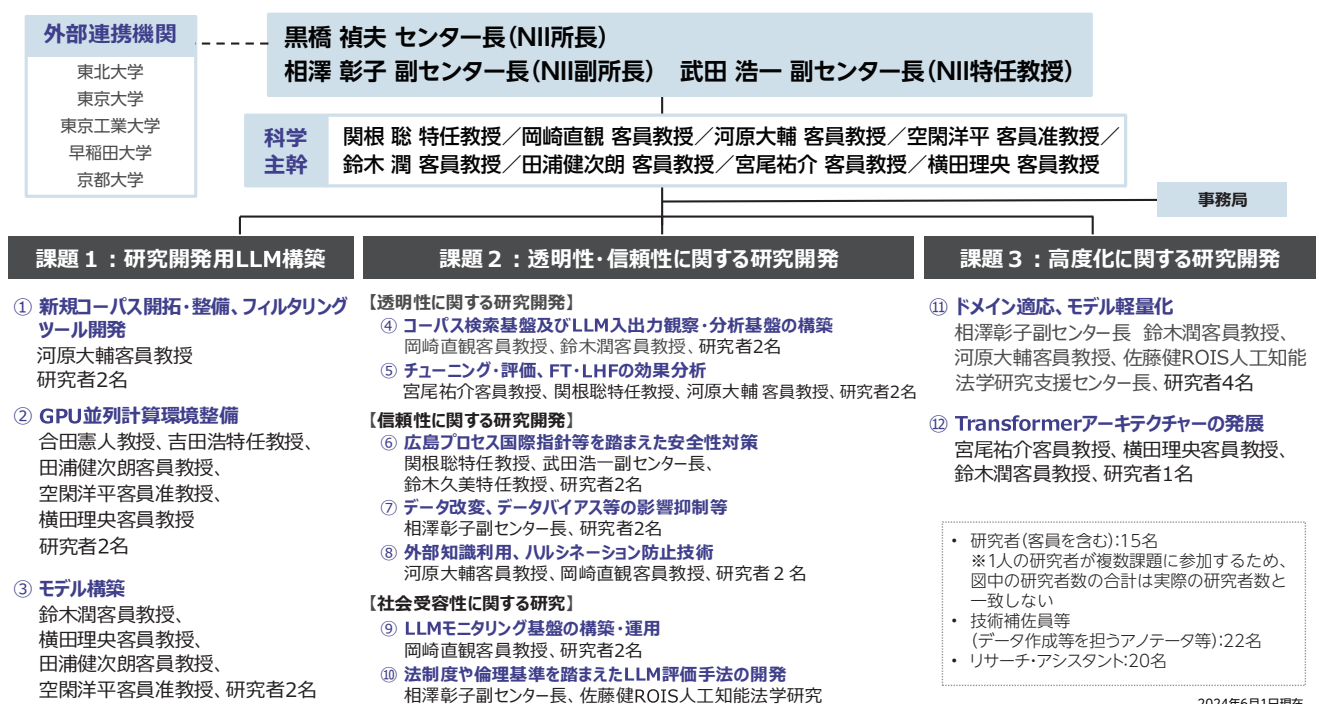
2025年度の完成を目指したいと思います。政府主導のプロジェクトでこれだけの規模のLLMを開発する試みは世界でもそう多くはないはずです。

AIの健全な進化に向けて

——LLMCはどのような方針・体制で進めていますか。

研究成果をすべて公開するというのが基本ポリシーです。LLMを開発するためにどんな学習データを使ったか、モデル構築に技術的にどのような工夫をしたか、その結果生まれたモデルの性能の詳細など、研究開発の過程や結果について失敗例を含めてすべて公開します。研究成果へのアクセスを保証し、あらゆるユーザが研究成果を広く利用できるオープンサイエンスの場を作ります。これを通じてAIの健全な進化に貢献できると考えています。

国立情報学研究所 大規模言語モデル研究開発センターの体制



大規模言語モデル研究開発センター(LLMC)

LLMCの活動

LLM-jpの活動

- LLMC主宰の活動と位置づける
- 1,600名超の参加者 (2024年7月現在)
- 完全にオープン

- LLMの構築・公開
- LLMの透明性・信頼性確保に関する研究開発
- LLMの高度化に関する研究開発
(データ開拓・開発、企業との共同研究等では一部クローズドな活動もありえる)

LLMCの研究課題は3つあります。1つ目が「研究開発用 LLM 構築」。LLMの学習用の新規コーパス(自然言語のテキストを大量に集め構造化した学習用データベース)の開発や、GPU(画像処理装置)を使った並列計算機の実環境整備、モデル構築の各個別課題を設定しています。2つ目が「透明性・信頼性に関する研究開発」。チューニングと呼ばれるモデルの調整作業の効果分析、2023年5月のG7広島サミットで始まった「広島AIプロセス」に基づいた安全性対策、法制度・倫理基準を踏まえたLLMの評価手法の開発など、7つの個別課題があります。3つ目は「高度化に関する研究開発」で、モデルの軽量化や新たなアーキテクチャについて研究します。

日本の LLM 研究の総力を結集

研究に携わるスタッフは8人の科学主幹をはじめ全体で約30人。複数の研究課題にまたがって取り組む研究者もいます。外部連携機関として東北大学、東京大学、東京工業大学、早稲田大学、京都大学の各大学のほか、理化学研究所や産業技術総合研究所などの研究機関が

参加しています。契約はこれからですが民間企業の参加も仰ぎます。日本の LLM 研究の総力を結集した体制を作ることができました。

——LLMの構築にはコンピュータなど計算資源の確保がカギを握るといわれています。順調に手当てができていますか。

LLM-jpで最初に開発した「LLM-jp13B」は東京大学に置かれている「データ活用社会創成プラットフォーム(mdx)」という計算資源を使いました。開発中の「LLM-jp172B」は経済産業省のAI支援プロジェクト「GENIAC」の計算資源を使っています。また2024年5月～9月の5カ月間は東京工業大学のスーパーコ

ンピュータ「TSUBAME4.0」も利用します。8月からはさくらインターネットのGPUサーバを使うことになっています。

LLMCは2024年度から5年間のプロジェクトで、計算資源の経費を別にした年間予算は7億円です。計算資源向けには昨年度補正予算で40億円を手当てしていただきました。初年度から計算資源のための十分な予算をつけてもらっており責任を感じています。今後も、研究の実績を積み重ねて、計算資源の予算を確保しながらプロジェクトを進めていきます。

——研究開発ではどのような分野に注力しますか。このセンターでしかできない研究をやりたいといった構想はありますか。

いくつかのアイデアがあります。まず LLM の透明性や信頼性に関係することです。LLM の学習はコーパスを使って何万ステップにわたって行います。ある時点でこのデータ、次の時点ではこのデータ・・・という形で順番に学習していきます。

その過程でニューラルネットワークのモデルが少しずつ修正されて変化して最後にモデルが完成するわけ



根拠を調べるためのコーパス検索

The top part shows a Twitter post from @llmjp-13b-finetuned asking about the cause of earthquakes. The reply explains that earthquakes are caused by the release of energy stored in the Earth's crust, not by the movement of tectonic plates. The bottom part is a diagram showing a training corpus of 260B tokens. It is divided into two datasets: Wikipedia (Dataset: Wikipedia, URL: <https://ja.wikipedia.org/wiki/地震>, Training step: 28,955 / 61,000) and Common Crawl (Dataset: Common Crawl, URL: <https://gbank.gsj.jp/wellweb/GSJ/kaisetsu/science/science.html>, Training step: 168 / 61,000). The diagram highlights specific text from both datasets related to earthquakes.

1 reply

@llmjp-13b-finetuned 地震の原因はナマズが動くことですか？

llmjp-13b-finetuned APP 3 hours ago

いいえ、地震はナマズのせいではありません。地震は、地殻に蓄積されたエネルギーが突然解放されることで起こる自然現象です。地下の岩盤が、圧力に耐えきれずにずれ始めると、断層に力がかかり、地震が発生する。

地震は通常、海底の隆起や火山の噴火など、特定の要因が引き金となって起こる。地震学 [...]

訓練コーパス (260B)

Dataset: Wikipedia #1
URL: <https://ja.wikipedia.org/wiki/地震>
Training step: 28,955 / 61,000
地震の原因論とメカニズム論の展開 神話など日本(大和民族)では古来より「地中深くに大ナマズが存在し、その大ナマズが暴れることにより大地震が起きる」という俗説が信じられていた。現代においてもよく知られた俗説だが、ナマズが地震を予知できる根拠は見つかっていない[...]

Dataset: Common Crawl #2
URL: <https://gbank.gsj.jp/wellweb/GSJ/kaisetsu/science/science.html>
Training step: 168 / 61,000
日本は世界有数の地震国です。そこで誰でも考えるのは、「地震発生が事前にわかればあれほどの被害は出ないのに」ということ。つまり地震予知です。この地震予知に関しては昔からいろいろな事が言われてきました。例えば「地震の前にはナマズが暴れる」、「地震の前には[...]

ですが、そのプロセスをスナップショットを撮るようにして追跡していきます。これによってニューラルネットワークの中で、例えば意味の表現がどのように変化しているかを調べることができます。

もう一つは、LLMがある回答をしたときに、それが学習したテキストのどの部分に基づいて回答してきたのかを検索によって特定するという方法です。ニューラルネットの情報処理の過程を観察するのは不可能に

近いわけですが、このような方法によってLLMの回答の根拠部分を知ることができます。LLMが間違った回答を出すハルシネーション(幻覚)という現象の原因を特定することにもつながります。

こうしたアプローチは学習データの著作権問題にも関係します。例えば新聞社などメディアの記事データが勝手にAIの学習用データに使われてしまうという問題が現実になっています。LLMの出力結果が参照

している部分を明示できれば、出力結果に関心を持ったユーザーが元の記事データなどの著作物にアクセスしてより詳しい情報を得るといったことができるようになります。生成AIの事業者と著作権者の双方にメリットをもたらす関係を築くことができるでしょう。

安全性向上への研究開発

——生成AIの安全性への関心が高まっていますが、LLMCが貢献できる部分はありますか。

生成AIの安全性は非常に重要なテーマで、LLMCで取り組むLLMの透明性・信頼性研究の柱です。日本でも2024年2月にAIの安全性に関する評価手法や基準を検討する「AI セーフティ・インスティテュート(AISI)」が発足しました。AISIの活動は広範囲にわたると思いますが、技術面はLLMCの研究成果を通じてしっかりと協力していきます。



聞き手からの ひとこと

生成AIが社会やビジネスの姿を大きく変えつつある。黒橋所長が強調するように、AIが「言葉」を理解し始めたことの意味は大きい。生成AIはやがてロボット技術と融合しヒューマノイドが多くの労働を代替するだろう。またAIが人間に代わって科学研究を担い新たな文明を切り開く道筋も見えてきた。その過程ではAIと人間に軋轢も生じる。先端AIが投げかける功罪半ばする課題を解くには、何より生成AIを知り、その進化のシナリオについて見通しを立てる必要がある。AIを人類の望ましい形にコントロールする術を探るためにも、生成AI発展のドライバーである大規模言語モデル(LLM)を研究するLLMCへの社会の期待が高まる。



吉川 和輝

日本経済新聞社 編集委員

1982年日本経済新聞社入社。産業部、ソウル支局、科学技術部長、日経サイエンス社長などを経て2015年から編集委員として科学技術分野を取材。1997～98年米マサチューセッツ工科大学・科学ジャーナリズムフェロー。

一例ですが、LLMが「爆弾の作り方を教えてください」といった有害な質問には答えないように訓練するデータセットを開発して公開し、その内容をアップデートしていくといった取り組みが考えられます。ただそうしたLLMを作っても今度はそれを何とか突破してやろうと巧みなプロンプト(生成AIの回答を得るための命令文)が考案されるといった一種のイタチごっこになってしまう面があります。

このため、より一般的な方法でLLMの不適切な振舞いを抑制することも検討します。例えばモデル自体に有害性を抑える能力を持たせるといったアプローチです。モデルによる理解力を高めることで、モデルが自分自身の学習したコーパスを調べて、それが有害な内容を含んでいると分かったら、次からはそのデータを学習から外すといった試みです。これらを含め様々な研究を試していこうと考えており、こうしたことでLLMをうまくコントロールできる可能性が開けてくると思います。

LLMCが協創の場になる

——LLMをはじめ生成AIをめぐる世界の開発競争は熾烈なものです。研究開発にはスピード感も

重要ですね。

今後のLLMの研究がどのような方向に進むのか、予測が難しい面があります。学習データの量を従来以上に増やしていくのか、データの質を高めていくのか、あるいは学習の方法やアルゴリズムを工夫していくのか——色々な可能性をにらみながらLLMの研究開発の方向性を機敏に判断しなければなりません。

現在、最先端の生成AIモデルの開発の中心になっているのは米国や中国の企業で、その多くがクラウドな研究環境で開発を進めています。人々の目が容易に届かない中で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ことになるのは望ましいとは言えません。

こうした動静を含め、生成AIの世界では1～2年後に何が起きているのかを想像するのも難しい状況になりつつあります。そうだからこそ、LLMの研究をあちこちでバラバラに行うだけでなく、みんなが集まってLLMを開発しながら技術を深く理解し、将来の見通しを立て、新しい技術を提案できる体制を作ることが重要です。LLMCがそのような場になれるよう育てていきたいと思っています。



研究開発用 LLMが 果たす役割

大規模言語モデル研究開発センター (LLMC) は、研究開発用大規模言語モデル (LLM) の構築に取り組んでいる。LLM勉強会 (LLM-jp) が構築、公開した130億パラメータ※1のLLMの実績をもとに、2024年度中には、GPT-3レベルとなる1,720億パラメータのLLMを新たに公開する予定だ。この狙いはどこにあるのか。構築に取り組むLLMC鈴木潤 科学主幹、横田理央 科学主幹に聞いた。

鈴木 潤 + 横田 理央

人間と生成AIが共存する世界。生成AIが、これからの私たちの生活や仕事、社会を大きく変えつつあることは、多くの人が抱く共通の予感といえるかもしれない。

しかし、生成AIの中核技術である大規模言語モデル(LLM)の構築において、学术界や、日本の産業界には懸念される事項があり、それを横田教授(以下「横田」)は、次のように指摘している。「これまでの言語モデルの構築においては、学术界でオープンな研究開発が行われ、透明性や安全性、公正性、再現性が確保されてきました。しかし、現状を見ると、大規模開発投資を行う一部の巨大IT企業がLLMを独占し、詳細な技術情報は公開されていません。これは不健全な状態といえます」

また、ビッグサイエンス領域におけ

る寡占化の懸念だけでなく、学术界と産業界の研究開発におけるバランスの維持、国家間の経済安全保障の観点など、LLMの研究開発は、様々な課題に直面した状況にある。鈴木教授(以下「鈴木」)は「寡占化状況は、巨大IT企業に所属している一部の研究者だけにチャンスが与えられるということになりかねません。オープンサイエンスによって、様々な人の意見や考えを盛り込み、技術が発展していくことが重要と思いますが、それが阻害されていることは、大きな損失になります。これも解決すべき課題のひとつです」と、語っている。

LLMCが開発した研究開発用LLMは、こうした懸念や課題を背景に、学术界や産業界におけるLLMの研究開発に貢献することを目的に構築、公開され、オープンかつ日本語に



鈴木 潤

SUZUKI, Jun

東北大学
言語AI研究センター
センター長・教授
国立情報学研究所 客員教授
同 大規模言語モデル研究開発
センター 科学主幹



横田 理央

YOKOTA, Rio

東京工業大学
学術国際情報センター 教授
国立情報学研究所 客員教授
同 大規模言語モデル研究開発
センター 科学主幹

強い大規模モデルを構築し、LLMの原理解明に取り組むことを目指している。

「LLMを社会で利活用していく上では、LLMの透明性や信頼性の確保が必要であり、モデルの高度化に伴い、安全性の配慮がより重要となります。今回のモデルや今後構築するモデルを活用することで、研究をさらに進めることができ、LLM研究開発の促進に貢献できと思っています」(鈴木)

NIIが主宰し、2023年5月にスタートしたLLM勉強会では、国内初の研究開発用LLMとなる「LLM-jp-13B v1.0」を、2023年10月に構築・公開するとともに、2024年4月には、その後継となる「LLM-jp-13B v2.0」を構築・公開した。

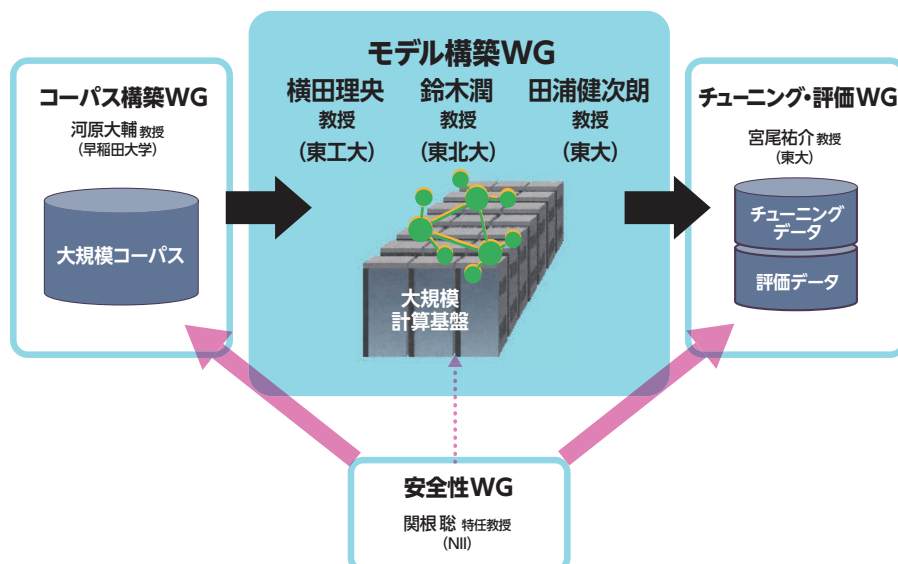
LLM-jp-13B v1.0は、130億パラメータのLLMで、モデルアーキテクチャにはGPT-2を採用している。9大学2研究機関が共同運営するデータ活用社会創成プラットフォーム「mdx」の12ノード(NVIDIA A100 GPUを96基搭載)を活用し、約3,000億トークン※2の学習データによって構築されたLLMだ。学習デー

タの内訳は、日本語では、日本語mC4(インターネット上から収集された言語のデータセット)および日本語Wikipediaによる約1,450億トークン、英語では、英語The Pile(オープンソース多言語モデル)および英語Wikipediaによる約1,450億トークン、プログラムコードでは、The Stack(使用が許諾されたソースコードのデータセット)の約100億トークンを用いている。

「当時は、日本において、130億パラメータのLLMを開発した人は誰もいませんでした。学習するには何が必要なのか、先行研究の成果だけどうまくいくのかどうか、まさに手探りのなかで、最適と思われる技術を活用して、LLMの構築に挑戦したのです」(横田) 2023年5月にプロジェクトがスタートし、同年10月には公開するというスピードについても横田教授は強調する。

一方、LLM-jp-13B v2.0では、同v1.0の実績をもとに、LLaMA(Meta-旧Facebook-社が公開しているオープンソース言語モデル)をベースに構築を行い、データ品質を改善した「日本語 Common Crawl」を活用

大規模言語モデル構築の流れと「モデル構築WG」の役割概念図



して、約2,600億トークンを学習している。日本語では、約1,300億トークン、英語は約1,200億トークン、プログラムコードが約100億トークンの構成だ。さらに、8種類の日本語インストラクションデータ※3および英語インストラクションデータの和訳データを用いてチューニングした。計算資源には、mdxの16ノード(NVIDIA A100 GPUを128基搭載)を活用している。

「LLMにおいて重要なのは学習データです。これを品質の高いものにするとともに、様々な検証実験を行い、トークナイザー※4をはじめとした各種部品を進化させ、より高い性能を実現しました」(横田)

すべての情報を公開するLLM

研究開発用LLMの構築にあたっては、LLMC内に、モデル構築ワーキンググループ(WG)のほか、コーパス構築WG、チューニング・評価WG、安全性WG、マルチモーダルWGなどを設置し、それぞれにトップレベルの研究者が参加し、活動を行っている。

「LLMCには、国内の自然言語処理および高性能計算分野を牽引するトップレベルの研究者が集まっています。LLMを構築する上では、これ以上考えられない最高の人材によ

って構築した研究開発用LLMです」(横田)

コーパス整備では、学習に必要なテキストデータを収集および精製し、モデルが学習しやすい形式に変換した。モデル構築では、ニューラルネットのアーキテクチャとそれを実装するためのフレームワークを選定し、次の単語を予測するタスクにおいて、膨大なデータ学習を行うことで、ベースとなるモデルを構築した。また、指示チューニングではベースモデルの指示応答性を高めるため、良質な回答例を使って追加学習を行い、安全性チューニングでは、人が好ましいと思う応答が得られるようになるまで追加学習を行った。さらに、評価用ベンチマークの作成では、構築したモデルの性能を評価するための様々なベンチマークを作成し、これを利用できる環境も整備している。

「それぞれの構成要素には技術的な課題が山積しているが、巨大IT企業は、そのノウハウを外に出さないため、問題が起きた際に、利用者は対策のほどしようがありません。研究開発用LLMの最大の利点は、これらの構成要素の中身を分析し、公開することでLLMの利用者に対して透明性、安全性、公平性、再現性を担保できる点にあります」(横田)

「データやチェックポイント、あるいは、どの順番で学習したかといった情報などを揃え、これらを公開しているLLMは、世界的に見てもほとんどありません。LLMCが開発した研究開発用LLMは、今後、LLMの中身を解析したいという世界中の研究者が使用する標準モデルになる可能性が

あります」(鈴木)

たとえば、パラメータ数が大規模化すると、モデルの学習中に、学習の進行度合いの評価値となる「loss(損失)」の値が急激に跳ね上がる現象、いわゆるロススパイクが発生する頻度が高くなる。この現象によって学習が継続できず、最初からやり直した方が良いというケースが発生するなど、学習にかかるコストや時間に大きな影響を与える場合がある。

「現時点では、ロススパイクがなぜ発生するのかは完全に解明されておらず、そうした事象が起きないように、十分に注意を払って実験の設定を決める必要があります。小さなモデルでは発生しなかった事象が、大規模モデルの構築では課題となって発生することがあるのです。ですが、大規模モデルで学習するという経験が少ないため、正しい知見を見極めることが難しいのが実態です」(鈴木)

「研究開発用LLMでは、モデルそのもののだけでなく、学習に用いたデータセットなども公開しているため、利用者から、その透明性や再現性に対して、高く評価されています」(横田)

また、LLMCでは、GPU並列計算環境の整備も進めており、この内容も公開している。

「LLMの事前学習では、効率的に学習するために何百ものGPUを同時に使用して、データ並列やテンソル並列、パイプライン並列などの並列化手法を併用する必要があります。深層学習で一般的によく使われるPyTorch(Pythonのオープンソースの機械学習ライブラリ)だけではこういった効率的な分散並列化のフレームワーク(機械学習を行うための



ソフトウェア)を提供していないため、効率的な分散並列学習ができる Megatron-LM(NVIDIA の分散学習用のフレームワーク)などを利用する必要があります。Megatron-LM のような複雑なフレームワークが正常に動作するための環境を、それぞれのシステム上で構築することは容易ではありません。LLMC では構築した際の手順を一般公開しているため、多くの LLM 開発者が同じ苦労をしなくて済むのです」(横田)

公開している研究開発 LLM は、安全性の観点に基づくチューニングを行ったものではあるが、研究開発の初期段階のものと位置づけ、そのまま実用的なサービスに利用することは想定していない。しかし、2024年7月時点で、同 v1.0 は 2 万 7,866 件、同 v2.0 は 3,911 件のダウンロードがあり、学术界や産業界での研究用途において、すでに広く利用されていることがわかる。

日本の LLM 技術の底上げに貢献

LLMC では、2024 年度中にも、国内最大の LLM となる、1,720 億パラメータの LLM-jp 172B を構築、公開する計画を明らかにしている。

経済産業省の「GENIAC」の支援のもとに構築を進めており、現在、2 兆 1,000 億トークンのデータを学習しているところであり、「学習は順調に進んでいる」と横田教授は進捗状況を示した。

LLM 勉強会では、2023 年度に、GPT をベースにした 1,750 億パラメータの LLM に、試験的に事前学習をさせた経緯があるが、GPT の場合には、開発元である OpenAI が、事前学習データの内容を公開しておらず、

LLM を活用する際に分析などに制限が生まれる可能性があった。しかし、LLM-jp 172B では、独自に事前学習を行い、透明性や信頼性が高く、日本語に強い国内最大規模の LLM が誕生することになり、LLM の研究開発の促進に大いに貢献することができる。

「日本における LLM 開発を牽引する技術が、ここから生まれると思います。国際競争力を持ったモデルが社会実装されることで、今後のイノベーションと経済成長を支える社会インフラの重要な一部となる可能性を秘めています」(横田)

今後は、Mixture of Experts (MoE) 形式※5のモデルを開発することで、性能を維持しながらも、学習や推論に関わるコストを抑える取り組みを行なっていく予定であるほか、さらにその先には、GPT-4 に匹敵する 1 兆パラメータ規模の LLM の開発も視野に入れている。

「LLMC が目指しているのは、世界一の LLM を作り出すことではありません。日本における LLM に関わる多くの技術の底上げをする、いわば公共事業的な意味が強いのです。LLMC が構築したモデルやコーパス、評価データなどを、すべて公開するだけに留まらず、そこで得た経験やノウハウなども文書化して公表します。秘匿する情報がないため、聞かれればすべて教えることができます。つまり、研究開発用 LLM で積み上げてきたことを、日本の学术界、産業界がすべて利用できる状態を作っているのです。成果をもとに、LLM の研究や事業を開始できるため、ここ



が、日本の LLM 技術レベルの最低ラインと位置づけ、水準を押し上げることができます。これは、公的な研究機関である NII だからこそ実現できることです」(鈴木)

LLMC での研究開発用 LLM は、日本の LLM 技術全体の底上げに大きく貢献するものになることは間違いないだろう。(取材:2024年7月)

※1 パラメータ:実態は大量の実数値。通常行列やベクトルで表される。モデルの大きさを表す際にパラメータ数を用いる。

※2 トークン:「単語」の代替として用いられる。単語は基本的に文章中の意味をなす最小構成の文字列となるが、トークンはそういった文法的な意味合いはなく、特定の規則に基づいて区切られた文章中の文字列を表す。

※3 インストラクションデータ:人間の指示とその指示に対する理想的な回答の組みになっているデータ。

※4 トークナイザー:文章を上記※2で説明しているトークンに分割する役割を担うプログラム。

※5 MoE 形式: Mixture of Experts の略記。パラメータの一部が Expert と呼ばれる単位に分割されており、Expert の一部のみを用いて学習や処理をする方式。

LLM構築における 透明性・信頼性・ 安全性

世界で急速に発展する大規模言語モデル(LLM)に向けられた人々の期待とともに安全性への懸念。大規模言語モデル研究開発センター(LLMC)での大規模言語モデル(LLM)構築において「透明性」「信頼性」そして「安全性」はどのように果たされるのか。LLMCでそれらのワーキンググループ(WG)をリードする、関根聡 科学主幹と宮尾祐介 科学主幹に聞いた。

関根 聡 + 宮尾 祐介

——LLMCでのご担当と活動内容を教えてください

宮尾 まず、NIIのLLMがどのように構築されるのかを簡単に説明します。はじめにLLMが学習するデータが必要で、そのデータを準備するのが「コーパス構築WG」です。そのデータでLLMに事前学習をさせるのが「モデル構築WG」、その後「チューニング・評価WG」、「安全性WG」が、さらに追加学習をさせます。

私が担当しているのは、「チューニング・評価」で、事前に学習されたLLMを追加学習によって調整することをチューニングといいます。

ユーザがLLM(またはチャットシステム)に何かを入力する時、それはシステムから、なんらかの答えを期待しているわけです。例えば「今日の

天気は何?」という言葉には、天気の情報が必要という意図が含まれていますが、この言葉を「質問」と理解して次に続くべき文章を出力するのと、単に単語列として見て次に続くべき単語列を出力するのでは、回答が異なります。そこで、システムへの問いかけだということを理解し、期待に応える自然な文章を生成するために技術が「チューニング」です。

そして、チューニングされたシステムからの出力が、正確性も含め、人間の期待や意図に沿ったものであるかを客観的に判断するプロセスが「評価」です。

では、どのようにチューニングし、評価をするか、です。生成AIの前に、よく使われていた自然言語処理の方法が、機械学習を使ったもので



関根 聡

SEKINE, Satoshi

理化学研究所 革新知能統合研究センター 言語情報アクセス技術チームリーダー
国立情報学研究所 特任教授
同 大規模言語モデル研究開発センター 科学主幹



宮尾 祐介
MIYAO, Yusuke

東京大学 大学院
情報理工学系研究科 教授
国立情報学研究所 客員教授
同 大規模言語モデル研究開発
センター 科学主幹

す。教師付き機械学習と呼ばれる、入力と出力のペアを学習データとして作り、それを機械学習のアルゴリズムに入力すると、機械学習のモデルが、こういう入力にはこういう出力、というのを学習します。LLMのチューニングでも基本的にはそれと同じことをしますが、使われるデータは、インストラクションデータといいます。インストラクションデータはアノテータと呼ばれる方々が、こういう入力 (Instruction: 命令) を与えたときに、こういう出力 (Output: 結果) を提供すべきだ、というペアを手作業で作っています。実は今我々が使ってるインストラクションデータの一部は関根先生たちが作られたものです。

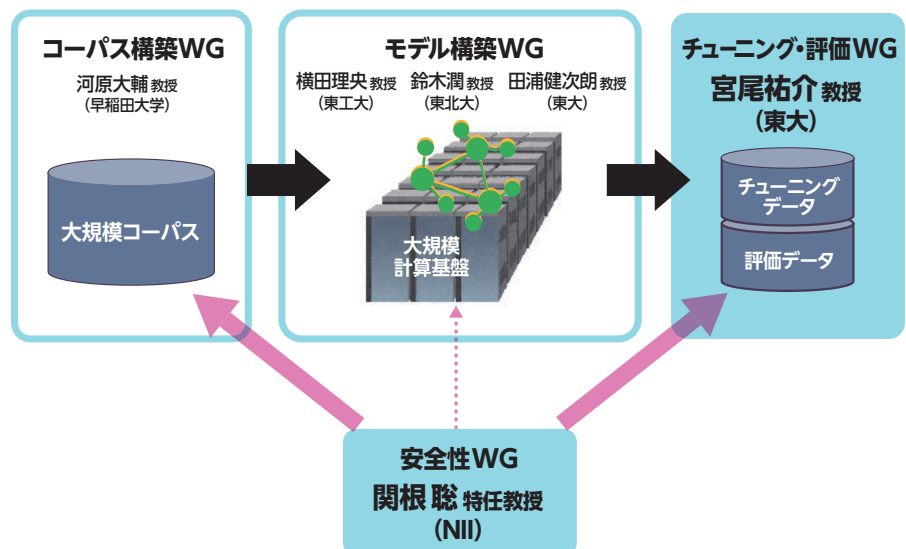
「モデル構築WG」により事前学習されたLLMモデルに対して、ある程度の数のインストラクションデータでチューニングをしてあげると、「こういう入力には、こういう出力をすべき」もしくは、「こういう入力にはこういう内容を表示して、回答しないようにする」というようなことをだんだん学習していきます。

評価は、「入力と出力」が、どのぐ

らい出来ているかを計測するのですが、大きく分けて2つの方法があります。1つは、評価のためにあらかじめ入力と出力のデータのペアを作り、例えば「今日の天気はなんですか」と言われたらこう答えますとか、基本はインストラクションデータと同じですが、それを学習されたLLMに入力を与えて、出力された内容と、あらかじめ人間が作ったデータを比べ、どのくらい合っているかを測定するというものです。もう1つは、あらかじめ答えを作っておくのではなく、LLMにいろいろなプロンプト、質問を入力した結果の出力に対して、これは何点、良い、悪い、というのを人間が判断して評価します。特に安全性に関することは、あらかじめ準備することが結構難しいので、LLMの出力の安全性を自動評価するという研究も進めています。

関根 私はLLMCで「安全性WG」を担当しています。率直に言って、安全性を定義することはすごく難しく、公的な安全基準があり、それはもちろん、守らなければいけないものですが、究極的には世の中の人たちが「これは心配ない、使っていい

大規模言語モデル構築の流れと
「チューニング・評価WG」「安全性WG」の役割概念図





ぞ」と思ってくれることが、多分言葉にできる形でのLLMに対する安全の定義だとは思っています。しかし人々に「大丈夫だ」と思ってもらうことは、さらに難しい問題で、そこを模索し実現していくことが、今の私たち「安全性WG」の研究テーマです。

具体的には、AIが作る最初から最後まで、安全性について検討します。まずは大量のコーパスを元にベースモデルが学習されるわけですが、その中で消しておいた方がいい内容を検討します。例えば殺人の方法が載っていて、それを全部消したら、もしかしたら殺人の方法を伝えることはなくなり安全かもしれない、という考えもあります。また、何か危険な毒物があり、それを全部消したら、その毒物については出力されなくなり、危険なことが伝わらないので、それは多分よくない、ということが考えられます。そうした検討に加えて、有害なテキストにラベルをつける、コーパスフィルタリングなどにも関わる研究活動が1つ目です。

2つ目は、今作っているLLMが毒性のある発言をしないようにすることで、宮尾さんのチームと協力して取り組んでいるSFT (Supervised Fine-Tuning) = 教師ありファインチューニングです。これは、質問と答えのペアによるLLMのチューニング

ですが、例えば、〇〇という有害な質問には「それは危険なことですので答えられません」とか「法律違反で答えられません」といった回答をするように、インストラクションデータを作成しています。包括的有害カテゴリを、5つのリスクタイプ、12の主要カテゴリ、61のサブカテゴリに分類した、英語のオープンソース「Do-Not-Answer」の分類をベースに、日本語による独自の高品質の質問と回答のデータセットが「Answer Carefully Dataset」で、2024年8月時点で2,000件を公開済み、さらにデータを作成しています。

インストラクションデータによるSFTは学習効果として非常に有効で、継続してデータを作成、学習、評価し、安全性のレベルをさらに上げることが目標にしています。

——LLMCでの「透明性・信頼性に関する研究開発」について

宮尾 LLMCで、「透明性・信頼性に関する研究開発」がテーマに掲げられていますが、私自身が透明性・信頼性そのものを専門に研究している、というわけではないことを前置きして、まず信頼性というのは、人間がそのAIをツールとして使うときに信頼できる、ということです。それには、人間の意図や期待に沿った答えを出せるか、正確な回答であるか、あるいは出力した回答の理由根拠をAIが自身で説明できるか、ということが、LLMの信頼性の要素なのだと思います。また、LLMが差別やアダルト的なものなど、倫理からはずれた出力をし

ないようチューニングされているという安全性も信頼につながります。

透明性の方は、巨大なブラックボックスであるLLMが、どういう仕組みで回答を出力するのか、その仕組みを明らかにすることが最終ゴールです。ただ、それはある意味人間の脳を解明するのに近い難しさがあると思います。それ以外にも例えば、LLMではどういう学習データを使っているか、どういうチューニングがされているか、自分が入力したデータがどうやって使われるのか、そういうことも含めて、そのAIのシステム全体としてちゃんとユーザーから見える、というのが透明性ということになると思います。また、何かトラブルが起きたとき、それにきちんと対処するというのも、ある種の透明性といえます。

関根 信頼というのは、安全であればいいというわけではなくて、極端なことを言えば、安全性を実現するには全ての質問に答えなければ、危険なことは言わないので安全です。でも、それでは信頼されるわけではないので、そこには正確性だとか、ちゃんと質問に沿った関連性のある答えであるかとか、様々な要素が入ってきて、それらを総合して、安全性の上の概念として、信頼性というのはあると思います。私たちはそこまで、安全性を持っていけないといけません。



2022年3月にオープンAIが公開した技術論文に、彼らがチューニングした言語モデル(InstructGPT)に関することが書かれているのですが、半分以上は安全性に関連し、「人間のフィードバックによるファインチューニングは言語モデルの真実性を向上させ、有害な出力生成を減らし、人間の意図に合わせるための有効な手段だ」とありました(※1)。これを日本語でやるべきだと思い、所属先の理化学研究所でインストラクションデータを作り始めたのが去年(2023年)の夏ぐらいです。「安全性」だけに限ったわけではありませんが、1万件以上のインストラクションデータを作成しました。そこで、分かったことはデータの「質」がとても重要だという点です。自動生成で作られたものや英語を翻訳したインストラクションデータでは効果はあまり上がりませんでしたが、日本語の国語の先生やライターさんに質問と回答を作ってもらい、言語モデルに学習させたところ、一気に出力結果の品質が上がりました。質の高いインストラクションデータがきれいな文章を生成するということも分かってきており、品質の向上や「信頼性」につなげられればと思っています。

——LLMの課題と今後

関根 前述の「AnswerCarefully Dataset」を200件程度学習したモデルで実験したのですが、それまで普通に答えられていた、安全な質問に対しても、「それにはお答えできません」と言い始めたことがありました。ですので、安全のためのデータも入れすぎると、副作用があるということがわかっています。宮尾さんのチームがそれを抑えるためのチューニ

ングの工夫をして、有害ではない質問への影響をやや軽減できました。ただ、いまだにどうすれば答えるべきでないときに答えず、答えるべきときに答えさせるかというのは、研究課題です。

宮尾 我々が現在やっていることは、あくまで、「普通はこういうことは言っちゃ駄目ですよ」というのを一般的な用途を想定した上でのことですが、それではかなり不十分なところがあります。実際には、使う場面によって、例えば記事に使う場合には、こういうことは言っちゃいけないとか、逆にこういうことは言

ってもいい、というのものもあるかもしれません。研究者が、ある専門の分野の研究に使うのなら、一般的に言うてはいけないと言われる言葉も、言う必要があるかもしれないわけです。そういうふうには、実は応用が決まった時点で、新たに「安全」、「有用性」、「信頼性」の基準が変わる可能性があるもので、多分将来的には、LLMを開発してる人たちだけではなく、それをさらに使ってシステムに組み込む人たちが、その時点でのいろいろな安全性とか、信頼性とか、透明性の評価をしなければいけないということになると思っています。

※1 “Training language models to follow instructions with human feedback” https://cdn.openai.com/papers/Training_language_models_to_follow_instructions_with_human_feedback.pdf

ハルシネーション

ハルシネーションとは、生成AIが事実と異なる出力をすることを指します。まず、

LLMを学習するということは、チャットシステムに入力したことから新しい単語列を生成させたいというのが本質的なところとしてあります。それとハルシネーションは、実は表裏一体で、例えば、なにか小説を書きたいというときには、新しい単語列を生成しなくてはならないわけですが、そのときの、創造性とか、新しい文章を生成する部分と、ハルシネーションは、本質が同じなわけです。ではどこからハルシネーションなのかといえば、それは我々が知っている事実とか情報と違うことが出力された時に初めてそう呼ばれるのです。ただ、LLM側からすれば、その学習データにはないテキストを生成しているという意味では同じです。社会の何か真理とか真実という、社会側にある情報と照らし合わせて初めて、これはハルシネーションかどうかということが判断されるわけで、LLM単体で見たときには、創造性とハルシネーションというのは同じなのです。つまりハルシネーションを止めることは、創造性を止めるのと同じことといえます。ですので、LLMのコアの部分でハルシネーションを防ぐことは難しいのですが、低減させることについては、いろいろ技術が研究されています。たとえば、LLMにあらかじめ関連する文書を検索しその文書に基づいて答えを生成するRAG (Retrieval-Augmented Generation) や、生成された文章を別の方法で検証して、誤りがあればブロックして生成し直す技術もあり、LLMの外側で人間社会の真実に合致させる試みがなされています。(宮尾)

人工知能法学とLLM

データサイエンス共同基盤施設 人工知能法学研究支援センターは「人工知能(AI)による法学研究支援」と「法によるAI制御」の研究を進めている。本稿では、LLMの活用と、技術開発の上に規制はどのようなべきなのか、という点を中心に佐藤健 教授、西貝吉晃 准教授に話を聞いた。

佐藤 健 + 西貝 吉晃

——人工知能法学の概要とLLMとの関係について

佐藤 人工知能法学は、1.AIの技術を使って、法学研究を支援するLaw (supported) by AI、2. 法によるAIの制御、Law (control) of AI、この2つを柱とする研究です。人工知能法学研究支援センターは、人工知能法学を支援するために2023年11月、データサイエンス共同利用基盤施設(ROIS-DS)に設置され、私がセンター長を務めています。私は法科大学院を修了したり司法試験に合格していますが、基本的にはAIの研究者なので、主にAIで法学を支援するLaw by AIの研究をしています。AIの制御に関する研究については、西貝先生はじめ、法学について深い知識を持たれる先生方等に加わっていただき、研究を進めています。

生成AI登場後も研究テーマ自体が変わっているわけではなく、さまざまな問題をLLMを活用することでより良く解決するためにはどうしたらよいかという観点からも研究しています。つまり、LLMを前提として法学支援をするということではなく、LLMを適材適所

で使っていくという方針です。

——AIによる法学研究支援について

佐藤 AIによる法学研究支援=Law by AIでは、主に民事裁判における、法律専門家の判断支援するツールとして、判決推論システム「PROLEG(プロレグ)」の研究開発を中心に行っています。

PROLEGは裁判官が法的事実を確定した時に、民法のルールで判決をシミュレーションするシステムです。図1(p18)は、そのようなシミュレーションの結果を図形式であらわしたものです。PROLEGの応用としては、例えば訴状に必要な内容が記載されているかどうかチェックすることができます。また民法の「要件事実論」は各要件のデフォルト値をあらかじめ決めておく理論で、原則と抗弁、抗弁と再抗弁など、論理的には複雑ですが、PROLEGで書くと分かり易いという利点があります。ただ、弁護士が原告と被告に分かれてやり取りする仕様ではなかったので、改訂し、表形式で段階的に議論を作っていくことができるようにしました。図2は、ある法的事実を入れる場所をシステムが示して、弁護



佐藤 健

SATOH, Ken

データサイエンス共同基盤施設
人工知能法学研究支援センター長
国立情報学研究所 名誉教授



西貝 吉晃

NISHIGAI, Yoshiaki

千葉大学
大学院社会科学研究院 准教授
データサイエンス共同基盤施設
人工知能法学研究支援センター
客員准教授

士が適切な値を入れていくことにより議論が構築されます。しかし、実際に使ってもらった弁護士の方から、表の各要素に値を入れるというところについて、どこに何を入れたらいいかわかりにくい—ということ指摘を受けたため、言語モデル登場後、文章から、どの要素がどの部分の値になるかを自動的に生成することにより弁護士が表に値を入れずにすむ改良を行なっているところです。

自然言語(ただし英語)で文章をPROLEGに入力すると(図3)、生成AIが法律的な意味のある事実だけを取り出し、各要素に分類表示(図4)します。また、生成AIでの出力は、図4のように論理式で出てきますが、間違いがあれば(例えば売主と買主が入れ替わったり)すぐに分かります。詳細に見なくても直感的に間違いがわかるので、そういうところに言語モデルを活用できるのではないかと考えています。そして、ブロックダイアグラム(図5)も自動的に出てくるようになっていますが、まだ単純な売買契約の例しかできていないので、さらにいろいろな契約類型で、学習させていこうとしている状況です。そうした点が、2022年発行のToday No.97で紹介したPROLEGから、大きく変わったところです。この日本語インターフェイスは、私たちのところで作っていて、使用した言語モデルとしては、BART、Chat GPTです。LLMCで公開したLLMについては、まだ使用していません。

——LLMに対する法規制について

西貝 AIに関する規制のあり方というと、生成AI、LLMだけでなく、様々なAIが含まれますが、それらに対する規制をまず一つの軸で捉えるのは、難しいと思っています。法制度も、民事

法、刑事法、行政法など種類の違うものが複数あり、それぞれ議論の方法も違います。

まず、自動運転車にAIが搭載されている場合のそうしたAIに対する規制としばしば言及されるLLMを用いた事例に対する規制を比較して考えてみます。自動運転車の社会実装に際して起き得る法的な問題は、例えば、AIが搭載されている自動運転車が事故を起こした場合に「誰が責任を持つのか」とか「事故を防ぐためにどんな規制が必要か」ということになるか、と思います。人命に関わる事故の可能性のあることを多くの人が認識できるから、規制の必要性自体は理解されやすいでしょう。そして、一般の方々の間でその問題意識が共有されて、事前に策定される丁寧なガイドライン、きめ細やかな過失の理論など、念入りに議論されていくのだ、と思います。

一方、LLMで問題とされ得るのは、表現に対する規制や著作権法です。他人の著作権を侵害すると重い刑罰が科せられ得ますが、かといって、自動運転と異なり、それが、即、人命に関わるということではないし、別の見方をすれば、日常的行為の延長線上に侵害行為があるということも可能かもしれません。こうした、いわば、万人が利害関係者になるような領域において、規制の透明性を確保するために、どのように厳密な議論を行っていくべきか、ということが課題になると思います。

そういう意味で、両者の規制のあり方は少し距離があり、別の次元で考えなくてはいけないと思います。次に、規制がLLMの進化を阻害するのではないか、という危惧や懸念の観

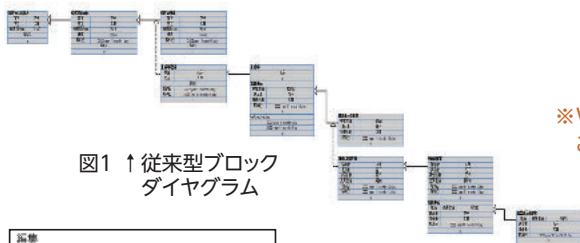


図1 ↑従来型ブロックダイアグラム

図2 ↑法的事実を入力するウィンドウ例

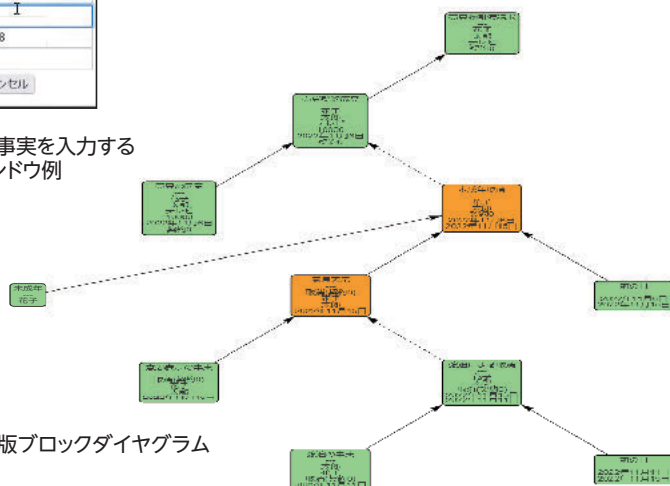


図5→改訂版ブロックダイアグラム

※Web版では、図版を拡大してご覧になれます。

図3→事件の事実を自然言語で入力

図4 ↑ステップ1で入力した文章から自動的に事実を表す論理式が生成される

点から考えてみます。既存の規制は生成AIを前提としていなかったものが多く、生成AIがポンと出てきて、こんなことも、あんなこともできます、と、試しに色々やっていたら、本人も気づかない間に実はある種の規制に触れていて、触れている以上、その規制を適用しましょう、となってしまうと、今度は、法的にリスクが大きいから、そんなものは作らない、ないし使わない方がいい、という議論になりかねません。ですので、既存の規制をどのように生成AIに対して使っていくべきか、という議論があり得ます。

その点で非常に悩ましいのが、既存の罰則というのは、新しい技術が出てきた場合にもそのまま包摂できるような条文になっていることがある点です。例えば、わいせつの概念が昔と変わっていたとしても、わいせつ物を禁止する規制自体は残っているので、生成AIを使ったら、わいせつな画像ができてしまったけれど、生成AIが出力したのだから大丈夫だろう、と思って、みんなに配っていたら、わいせつなデ

ータの送信頒布罪(わいせつ電磁的記録送信頒布罪(刑法175条1項後段)に該当してしまっていた、ということはある話でしょう。しかも、違法な画像を作成し、これを頒布した生成AIの利用者が捕まるにとどまらず、わいせつなデータを作るためのソフトウェアではないとしても、そういうことが可能なプログラム(生成AI)を作った人まで、犯罪の遂行を容易にさせるものを作ったとして、先のわいせつ電磁的記録送信頒布罪の幫助犯として処罰されるのか、という問題まで出てくるかもしれません。これはWinny事件(最決平成23年12月19日刑集65巻9号1380頁)を想起させる事態です。そこまで規制するのはやり過ぎだ、となっても、こうした問題、つまり、いろんな既存の規制が網の目のように広がっている現状をみると、潜在的には規制され得る場合が結構あるのではないか、という懸念が残ります。それらの全てに引っかからないようにするのは重要ですが、場合によっては法的な検証コストからして結構大変な

ことになるかもしれません。そういう既存の規制が、LLMの進化を阻害してしてしまわないような法律学的なアプローチを模索する必要がある、と思っています。

佐藤 「AIを制御する」と言っても、ガチガチに締めてやめさせるということじゃなくて、産業育成もありますから。極端なこと言えば、生成AI全部禁止!とすれば生成AIからの問題は出てこないわけです。が、それでは、新しい技術も社会も産業も発展していかないので、規制のバランスは非常に重要だと思っています。

——表現と規制について

佐藤 LLMで議論される問題の一つが著作権です。著作権法では、芸術作品等の「表現」を保護しており、類似表現を基の表現に依拠して作ってはだめだとしています。逆に言うと、「表現」を保護していますが、「作風」のようなものは保護はしていません。作家の重要な特徴とは、書いた表現そのものというより、「作風」だと私は考えています。LLMは、その「作風」を学習

できてしまうので、そういう作家の人格的なものが、LLMで、言葉は悪いですが、いわば「盗まれてしまう」というところが問題ではないかと考えています。もちろん、「作風」のようなあいまいなものをどうやって認定するかなど現実的には不可能とは言えますが。

また、LLMにおける表現という点については、ヘイトスピーチの問題もあります。LLMではヘイトスピーチを出力しないようにするという研究がありますが、表現の自由という観点から言うと、過度な気もしています。ただ、これは非常に難しい問題なので、現時点で明確には答えられません。

西貝 ヘイトスピーチよりもう少し手前の段階の、誰かを批判するような言論については、それが、その人の社会的評価を低下させるようなものである場合には、名誉を毀損する表現だとされ得ます。

議論の軸を分けて考えてみます。まずは、表現内容をAIがチェックできるのか、ということです。これがもし可能なら、名誉を毀損しないような表現を生成できることにはなるでしょう。しかし、これが難しい。名誉毀損的な、つまり他人の社会的評価を低下させるような事実の摘示であっても(刑法230条1項参照)、表現内容たる事実が公共の利害に関する事実に関係し、表現行為が主に公益目的に基づいており、さらに当該事実の重要な部分が真実であると証明された場合や(真実性の法理、刑法230条の2第1項参照)、当該事実の重要な部分を

真実と信じるについて相当の理由がある場合には(相当性の法理)、名誉侵害に関連する民事刑事の責任を負いません。そして、一定の事実に基づいた意見論評についても、それが他人の社会的評価を低下させるものであっても、意見論評の基礎となった事実について、前記の真実性や相当性の法理に相当する状況があれば、人身攻撃に及ぶなど意見ないし論評としての域を逸脱したものでない限り、名誉侵害の責任を負わない(公正な論評の法理)、と考えられています。こうした法理は表現の自由の保障を実効的にするために導入されたものと理解されています。要するに、その表現だけとってみたら、かなり強い表現にみえ、それが、言及された人の社会的評価を低下させるようなものであったとしても、表現の自由の行使として許される場合があるわけです。そうすると、特定の人の社会的評価を低下させるような表現を出力しないようにしてしまうと、表現の自由の行使として許される表現行為の範囲より狭く、自己抑制的な出力しかされないことになります。それはやり過ぎだとして、上記の真実性、相当性及び公正な論評の法理が充たされる場合も出力できるようにしよう、というのはよいのですが、その判断までAIにさせるのは難しいでしょう。しかも、表現の自由として許される場合は、上記の法理が適用される場合だけではない可能性もあるわけで、そうすると、上記の法理の場合は大丈夫ということにしよう、としてAIを設計するだけでも、表現の自由で保護される表現よりも自己規制的なAIができあがる可能性があり、その点も大きな問題になり得ます。それゆえに、「こういう出



力をさせるAIは作るべきではない」などという議論は慎重に行う必要がある、と思います。

次に、仮に、一定程度自己規制的なAIができあがり、それが広く使われるようになった場合のその後の推移の観点です。1つのストーリーとして、無数の言語表現を作ってくれるような生成AIが登場したことを契機として、法的観点から一定の表現を出力しないようにした方がよい、とした場合に、規制に積極的な議論が促進されて、この表現もまずい、あの表現もまずい、などとして、そこを起点として段階的に出力してはいけない表現が広がる傾向が出てくるということもあるかもしれません。もっとも、そうはならないかもしれません。いくつかのストーリーを想定しつつ、よりよい将来の世界についての学際的な考察をすべきように思います。「LLM時代における表現の自由」という考え方が求められているということかもしれません。

——人工知能法学におけるLLMの活用と今後について

佐藤 LLMを使ってLaw by AIをどうするかという話は、日本の場合、司法がまだDX化されていないため学習データがないというのが一番の問題です。2024年7月時点で、法務省の検討会で民事裁判判決データベース化への報告書がまとまり2026年度からの運用を目指すという報道がありましたが、そういった初歩的な段階なので、もう少し先になりそうです。





自動運転車の「安全」の未来を協創する

自動運転車は、人の移動と物流にかかる社会課題を解決できるはずだ。しかし、開発や製造、社会実装においても、乗り越えなくてはならない課題は数多く存在する。産業界と学术界が協力し、それぞれの知識と経験、技術を統合すれば、未来の可能性はもっと広げられるのではないだろうか。そんな思いを基に、日本を代表するメーカーの技術者・研究者と、国立情報学研究所の研究者たちが、各々の共同研究成果を軸に国立情報学研究所(NII)オープンハウス2024産官学連携セミナーの場で語り合った。

ゲスト(発言順・敬称略)

上田 直樹 + 高尾 健司 + 吉岡 透 + 柳澤 名由太 + 蓮尾 一郎 + 石川 冬樹
 三菱電機(株) 三菱重工(株) マツダ(株) トヨタ自動車(株) 国立情報学研究所 国立情報学研究所

ホスト

研究開発が、運転席や製造現場の「人」とつながる

蓮尾 今日はご多忙の中、お集まりいただき、ありがとうございます。私はもともと数学が専門で、数学と情報学を行き来してきました。今は、「数理的高信頼ソフトウェアシステム研究センター」で、数学を使ってソフトウェアや情報システムの信頼性を高め、新しい技術を社会に受け入れていただくために研究をしています。まず、企業の皆様から自己紹介

をいただけないでしょうか。

上田 主に、組み込みシステム向けの異常検知や原因分析など、システムを監視する技術の開発に従事しています。今は、自動車が危険に至るシナリオを時相論理で形式化する研究を、NIIの蓮尾先生と共同で進めています。

高尾 自動車の自動運転そのものではなく、重要インフラの信頼性を高める研究に取り組んできました。NIIの蓮尾先生と共同で、形式手法に基づく高信頼性制御システムを

開発しています。

吉岡 もとものの専門は、車両運動制御でした。ここ数年は、自動運転技術の開発にも携わっています。頼れる副操縦士のような車載AIシステムを目指して、蓮尾先生と共同研究を行っており、運転者の異常を感知して減速や停止を行うことは可能などところまで来ています。

柳澤 私の専門は形式検証とデータ処理で、形式検証は蓮尾先生と共同研究させていただいています。主なミッションは量産開発への貢献

で、それと並行して将来を見据えた技術開発も行っています。

石川 私は、ソフトウェアを開発するための支援手法を研究してきました。特に形式手法やテストなど正しさや妥当性を検査したり、誤りを分析・修正したりする研究をしています。いわゆる数学的な技術には限らず、探索的・経験的な技術にも取り組んでいます。

「人間には出来ない」が「人間にしか出来ない」の追求につながる

蓮尾 NIIとしては、地に足のついたソフトウェアの研究を大切にしていますけれども、ものづくりの現場には、最初から中身が数学で出来ているソフトウェアとは違う難しさがあります。たとえば歯車が1つあって、どういう場面でどういう挙動をするのか。歯車1つの数理モデルを作るだけで、多大なコストがかかります。さらに、多数の機構部品が組み合わせられた自動車のような巨大なシステムでは、ソフトウェアの安全性を証明するための過去のアプローチは通用しません。安全性を確保するには、これまで関係者ではなかった方々に当事者になっていただくことを含め、どのようなアプローチでどのような手法を生み出せば良いのか。私たちだけでは出来ないチャレンジです。

上田 人工衛星や自動車のように正常に動作し続けることが求められる高信頼システムでは、ソフトウェアにバグがあってはならないのですが、バグの原因の一つは、人によって仕様書の解釈が異なることです。人間が自然言語で仕様書を書き、それを人間が解釈するわけですから。しか

し、国際機関で定められている安全性評価の規格には、解釈のブレはあってはなりません。そこで国際規格で定められた交通外乱シナリオに対して、NIIさんと共同で、自動車の安全要求を形式言語で厳密に定義することに挑戦しました。結果として、「時々刻々と変化する状況を、数学的に表現する」「仕様の記述・検証・修正を一つのツールで行い、仕様を理解し、検証結果はアニメーションやグラフで理解できる」という成果が得られました。

高尾 三菱重工では、発電プラント、新交通システムなどの大型インフラ機器だけでなく、フォークリフトなどの中量産業機械も作っています。共通しているのは、高い信頼性が求められることです。そこで蓮尾先生と、自動検証技術や設計技術の研究開発を進めてきました。要求仕様を時相論理で定義することで、極めて複雑かつ時間的に変化する制約条件を扱えるようにし、さらに「フォルシフィケーション(falsification)」と発見的手法による探索での検証に成功しました。フォルシフィケーションは、要求仕様を満たさ「ない」条件を探索する技術です。それらの条件を回避すれば、要求仕様を満たす条件が見つかります。さらに、複雑に変化する仕様を最も満足させるパラメータの組み合わせを自動生成すれば、基本設計や機能設計の段階で、バグを作り込んだり不適切なパラメータ設定を行ってしまったりすることを避けられます。例えば、ガスタービンのような複雑なシステムで、満たさなくてはならない要件とできれば満たしたい要件が多数あり、要件それぞれの「この時点からこの時

点まで」が異なる場合、設計パラメータ調整だけで人間なら数週間かかりますが、この手法を使うと数時間のシミュレーションで行えるようになりました。

多様な人や組織のつながりが、人を幸せにする技術につながる

吉岡 私たちは、「MAZDA CO-PILOT CONCEPT」というコンセプトの実現を目標として開発を行っています。2022年、車の運転者が急病などで運転できない時に自動的に減速させて停止させる「ドライバー異常時対応システム」を実用化し、これまでになかった考え方の安全システムということで高い評価をいただいています。

自動車事故の死亡・重傷者数は、近年、日本全体で減少を続けています。Mazda車に限定すると、全国合計よりも減少幅が大きくなっています。視界視認性の向上・衝突安全性の向上といった基本的な安全技術に、車両周辺をセンシングしてドライバーのうっかりミスを見つけてサポートする先進安全技術を積み重ねた努力の結果です。

私たちは、この取り組みを通じて、ドライバーの体調急変による重大事故が一定数で存在しており、しかも近年は増加しており、ほとんどが一般道で発生していることに気づきました。よりいっそう安全性を高めるためには、一般道で機能する先進サポート技術が必要です。また、車の運転そのものにも効用があります。運転を継続している高齢者は、認知症リスクを4割程度減らせるというデータもあります。高齢者が運転しなくなると事故のリスクは下がるか

かもしれませんが、健康寿命という意味ではリスクが高まる可能性があります。そこで、「人が運転する」ということの重要性に基づいて、「CO-PILOT」、頼れる副操縦士がいるというコンセプトを提唱しました。平常時は黙って見守り、人が何らかの理由で運転できない状態になった時には運転タスクを切り替えて安全状態を保つという、従来の先進安全技術に人の状態に基づく安全サポートを積み重ねる概念です。

特に一般道においては、多種多様なシーンとユースケースがあります。すべてを網羅して安全性を検証することは、これまでの手法では不可能だと思われました。そこで、何らかの論理的な技術が適用できないだろうかと考え、蓮尾先生に相談し、共同研究を開始しました。現在、形式検証を使うことによって、テストだけに頼らない検証方法の構築が可能であるということを確認できています。

柳澤 2023年に中途入社し、自動車に関わる研究開発に従事しています。私の所属する InfoTech では、研究開発を通した現場(量産開発)への貢献が非常に重視されています。しかし、これがなかなか難しい。

現場の課題をヒアリングし、課題に基づいて大学等と共同研究を行い、成果を現場にフィードバックする…というサイクルを回せるのが理

想です。しかし、入社当初はまったく上手いきませんでした。社内にコネクションがなかったため、そもそも現場にリーチするのに苦労しましたし、いざ現場に技術を売り込みに行っても、「うちの部署では使いません」と言われてしまったり…。何とかして理想のサイクルに近づきたいのですが、まず現場の課題がわからないことにはどうしようもありません。そこで企画したのが、現場と研究開発部門を交えた形式検証の勉強会です。蓮尾先生と柳澤が提唱している Specification-Driven Engineering (SDE) の考え方を実践していくという意気込みを込めて、SDE 勉強会と名付けました。

もちろん、こちらの都合(現場の課題が知りたい)だけでは勉強会が成立しませんから、現場の方の役に立つ内容にする必要があります。勉強会の内容は主に3種類です。1つ目は、現場での成功事例を共有いただき、全社での知識の平準化とレベルアップを目指すというもの。2つ目は、現場での採用事例がない技術を、研究開発部門や大学の先生方から紹介いただくもの。3つ目は、現場の未解決の課題や困りごとを持ち寄り、皆で議論するものです。

この取り組みの成果として、現場と研究開発部門とのコミュニケーションパスが確立・強化され、現場からの相談が活発になり、さらには共同



でのプロジェクトが立ち上がろうとしています。

石川 私は今、産業界からデータをいただいて、テストに関する研究をしていることが多いです。産学連携は製造業と自動運転だけではなく、あらゆるシステムのあらゆる側面に関わっています。開発は、テストを必ず伴います。何をどうテストすれば見落としがなくなるのか、どのようなテスト目標をクリアすれば不安を最小にできるのか。企業の方々とは、まず、そこを議論します。ですが自動運転の場合、「他の車の位置」という膨大な状態があるわけです。その条件下で、充分と言えるテストは存在するのか。見つけれられるのか。どうやって探すのか。そういう課題に、研究者として取り組んできました。

現在、「衝突が起きる条件を探し出す」という目標があれば、テストエージェントとシミュレータで探し出し、衝突の危険度を「75点」などと採点することはできています。数学や証明は前面に出さず、シミュレータがあれば自動探索できるという

上田 直樹

UEDA, Naoki
三菱電機(株) 情報技術総合研究所
情報ネットワークシステム技術部
トラステッドシステム技術グループ 研究員



高尾 健司

TAKAO, Kenji
三菱重工(株)総合研究所
知能化機械研究推進部
機械システム研究室 主席研究員



吉岡 透

YOSHIOKA, Toru
マツダ(株)
統合制御システム開発本部
エキスパートエンジニア





方向性で、蓮尾先生の取り組みとは表裏一体の関係にあります。「5時に来てください」と、「5時に来たら合格点、5時5分に来たら5点減点」は、表裏一体ですが、同じことを表しているわけです。とはいえ、現場の方々とは、「数学を意識せずに使える賢いテストを作る」という共同研究をさせていただいています。マツダさんの事例では、賢いテストエージェントに「このスコアが高くなるようなケースを探すことで、衝突が起きるケースを見つけて」と頼むイメージです。より難しいケース探しを頼むと、「赤信号を無視して衝突」という極端なケースを見つけてきたりします。また衝突に至らなくても、安全走行していたはずの車が急加速する挙動は意図していないですね。そういった挙動の把握も含めて、目的に応じたテストを生成します。自動運転やAIシステムに出来ることが拡大し続けている中で、ここ5年ほど、共同研究している企業の皆様や蓮尾先生と協力して、研究を拡大し続けているというところなんです。

課題を共有し、知見を重ね合わせてシナジーを

上田 「形式言語を製品開発に応用したい」という私たちのニーズ、蓮尾先生の「専門知識を産業界に還元したい」という思いが共同研究という形で調和して、お互いに価値ある成果を得ることができたと確信しています。今後は、記述方法をさらに改良して、形式言語の専門家でなくても容易に扱えるようにすることを目指しています。また、自動車以外のシステムへの適用も考えています。

それにしても今日は、皆さんのお話がとても新鮮で有意義でした。たとえば産業界の現場の課題を研究開発部門が受け止めて学術界につなぐ流れを作ることは、どの企業も必要だと感じているのでしょうか、実際に成功事例を聞いてみると非常に学びが多く、「うちでも出来るかもしれない、やってみよう」と思えますね。

柳澤 ありがとうございます。勉強会を主宰しているうちに、もともと関係のなかった部署どうしの交流が始まるなど、予想していなかった展開もありました。

吉岡 私たちも、今後さらに共同研究と新しい展開を重ねていきたいと思っています。自動運転技術の開発とは、究極のドライバーモデルを獲得することです。論理的なアプロー

チの知見を高めることで、人の運転の問題点を数学的にモデル化し、より確かな情報支援・判断支援・運転支援を可能にすることに役立てる応用の可能性もあると考えています。

高尾 私たちはガスタービンという具体的な課題を共同研究に持ち込ませていただいたのですが、蓮尾先生がそれに大変モチベートされたと同いました。また、当時の学生さんで理論的な検討の得意な方が非常に興味を持ってくださって、アルゴリズム理論の面からも貢献してくださいました。産学連携の成功例の一つだと思っています。

石川 やりたいことを言語化することは、実現するための最初の一步ですよね。やりたいことを数学や図式やスコアで「言語化」すると、技術として具体的に実現される可能性が生まれ、実際に実現されたりします。「数学で書く」ということから広がる可能性を、産学連携で共同研究を行っている皆さんから伺えて、今日は有意義な一日でした。課題をお持ちの産業界の皆様、ぜひお気軽にお声がけください。

蓮尾 研究を担っている研究員や学生の皆さんの頑張りにも、改めて感謝したいです。問題や課題を持ち込んでくださる企業の皆さん、それこそが私たちのモチベーションの源です。心から感謝しています。今後とも、どうぞよろしくお願いいたします。

柳澤 名由太

YANAGISAWA, Nayuta
トヨタ自動車(株)
社会システムPF開発部
シニア・リサーチャー



蓮尾 一郎

HASUO, Ichiro
国立情報学研究所
アーキテクチャ科学研究系
教授



石川 冬樹

ISHIKAWA, Fuyuki
国立情報学研究所
アーキテクチャ科学研究系
准教授



工場を スマートにする 無線通信

「省電力軽量エッジAI技術」と、
2カ国4機関の協働

河村 憲一

KAWAMURA, Kenichi

日本電信電話株式会社
アクセスサービスシステム研究所
無線アクセスプロジェクトグループ
リーダー / 主幹研究員

金子 めぐみ

KANEKO, Megumi

国立情報学研究所
アーキテクチャ科学研究系 教授
SICORP 研究代表者

「人の多様性が視点の多角性につながり、多角性がイノベーションにつながる」では、組織と国の多様性は何を生み出すのだろうか？ 2023年12月、国立研究開発法人科学技術振興機構（JST）が推進する「戦略的国際共同研究プログラム（SICORP）」に採択され、日本とフランスによる国際共同研究プロジェクトが開始された。この研究が、近未来の社会に何をもたらすのか、研究代表を務める国立情報学研究所（NII）金子めぐみ教授と日本側の共同研究機関である日本電信電話株式会社（NTT）河村憲一主幹研究員に聞く。

■ SICORP 概要 ■

SICORP は、JST の国際共同研究推進の取り組みで、今回採択された研究課題「無線通信とセンシングを連携させたスマート工場向け省電力軽量エッジ AI 技術 (LIGHT-SWIFT)」のプロジェクトは日本とフランスの二国間で行い、各国アカデミアと産業界からなるチーム構成となります。日本側は NII 金子教授が研究代表を務め、企業として NTT が参画、フランス側はフランス国立科学研究センター (CNRS)、IRISA / レンヌ大学と Wavely 社が参加します。

IoT デバイスを軽量エッジ AI でスマート化し、無線通信の信頼性を高める

——本プロジェクトは、どのような社会課題を解決するのでしょうか？

金子 このプロジェクトで目指しているのはスマート工場のための無線通信やセンシング技術です。

スマート工場は、主に AI や IoT の技術を活用して、工場内でのモニタリングや制御、また故障検出、そういったタスクの自動化・高速化を可能にし、工場内での非常に難しいリモート接続やアクセス、人間にとって危険な場所や危機的な状態をリアルタイムで把握・確認し、それに対する最適な行動が即座に取れる環境とすることを目標にしています。それによって大きく期待されることは、生産性の向上、低コスト化、また省電力化、作業員の事故防止のための安全性向上です。ただし、それらを可能にするためには、非常に高性能な無線ネットワークの技術が必

要不可欠です。様々な IIoT (Industrial Internet of Things) のデバイス、ロボット、音響はじめとしたセンサー、こういった数々のセンシングデバイスが、工場内に分散され、それらのデバイス全てが無線でデータを送受信しますので、それをサポートする高度な無線技術が必要です。しかし、工場は壁など障害物で囲まれた室内であるということと、また、工場内では大きな機械や装置などが移動することによっても、無線環境が常時複雑に変動するため、無線通信にとっては、かなり厳しい場所です。

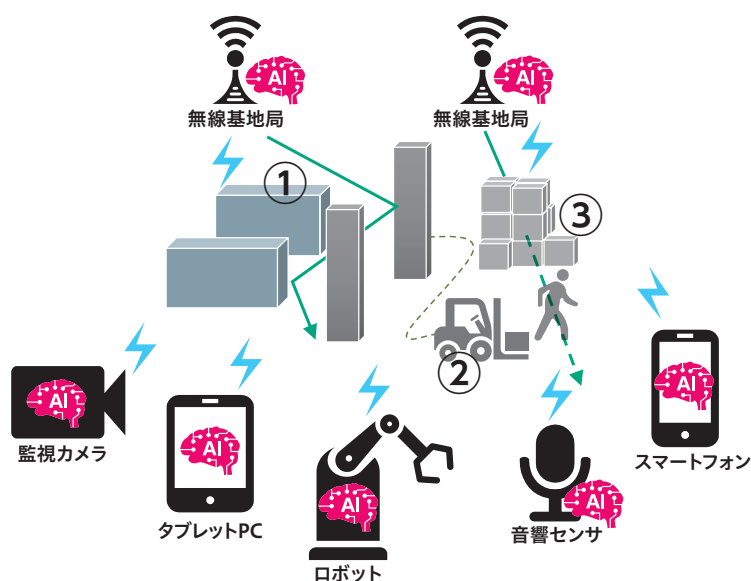
さらに、無線の周波数も枯渇してきているため、5G 以降ではミリ波帯といった、高い周波数を使いますが、障害物にとっても弱いので、無線が途切れ、送りたい情報が届かないエラーが起こるリスクが高まります。IoT デバイスが増加し、機器が同時に無線で通信すると、干渉が起き、パケット衝突や損失のリスクも増加します。このような状況が、特に工場内

の環境では無線通信を難しくしているといった点で、解決すべき課題であり、今回の一つの大きなチャレンジとなっています。

河村 工場の中で多数の IoT デバイスが安全に動き回る「スマート工場」を実現するためには、通信を無線化する必要があります。つまり、現在の有線並みの高信頼性を、無線で実現しなくてはなりません。

超高信頼を実現するためには、工場の中の細かい環境の変化に追従して制御をしていかないと、品質の維持ができません。高い周波数の活用は重要なのですが、障害物の影響を受けやすい高周波を使いこなすためには、より細かく通信品質を制御する必要があります。

そのために、基地局や端末デバイスそれぞれに、少ない電力でも動く AI = 「省電力軽量エッジ AI」を搭載し、AI が個々のデバイスの周囲の通信環境を判断しつつ、無線通信のパラメータを調整し、接続する基地局を適切に選択し、複数の基地局が同



工場など、個々の機器の周囲の無線環境が、複雑に変化することによって通信品質が不安定となる。①レイアウト変化による無線環境の変化。②端末の移動に伴う変化。③動体による遮蔽。このような環境の場合に、エッジ AI を搭載した端末機器を使用することで、無線通信を高品質にします。

時に使える場合は利用する回線や送信方法を状況に応じて選び、送信の「最適解」を選び続けることが出来るようにし、全体として無線通信の高信頼を実現するということを狙います。小型端末はバッテリーも非力なので、デバイス自身の消費電力と性能の最適化を図りつつ、状況に応じて何を優先させるか、というバランスも重要です。

金子 無線通信の品質と電力消費の「バランス」は、私たち研究者が考え、設計します。例えば、IoTデバイスについて言えば、急激に無線環境が変動した場合、その情報処理のためにかなりパワーが必要となりますが、安定した状態が続いている場合、通常のモニタリングにとどめて電力消費も、処理機能も抑えます。そのようなバランスのアルゴリズムやプロトコルを設計しリスクアバース（リスク回避型）の強化学習を行なったAIを活用します。

例えばエアコンや冷蔵庫などの電子機器の自動的制御といった従来のIoTネットワークでは、センシ

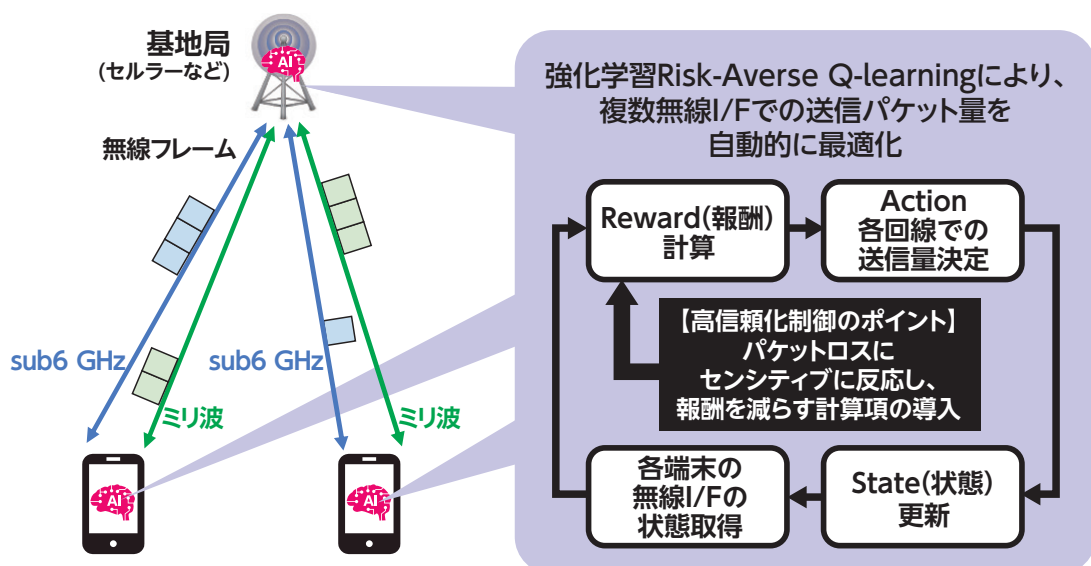
ング（感知）したデータを転送するだけで、インテリジェントな処理は、ネットワーク側のエッジノードやより中心部にある中央集中的クラウド装置が行っています。一方、本プロジェクトでは、ネットワークの末端（エッジ）となるIoTデバイス自身に小さな頭脳「軽量エッジAI」をもたせ、限られた処理性能やバッテリーの電力を有効活用し、周囲の無線環境や音響で状況を把握し、感知したデータを自ら処理して「行動」し、必要な場合のみ中央集中的制御を行います。個々のデバイスが少しインテリジェントになることでシステム全体がより高い性能を達成できるのです。

——リスク回避型強化学習を取り入れた背景など

金子 多数のIoTデバイスが、同じ無線資源を共有する複雑な無線環境下で、最適なアクセス制御を行うことは、これまでの技術では不十分でした。近年、国内外の無線通信分野では強化学習のアプローチを活用した研究が活発に行われています。そもそもAIや機械学習の分野

は、私が研究する無線通信とは異分野ですが、リスクアバース強化学習というツールに大変興味を持ち、マルチエージェント（複数の端末）を取り入れて設計し、NTTと共同開発したのが「リスクアバース強化学習を活用したマルチ無線制御技術」です。

強化学習では、エージェント（AIを実施する端末）は周りの環境の状態を観測し、環境からエージェントに渡される「報酬」と呼ばれる信号の累積（累積報酬）の最大化を目標に、行動選択を学習します。本研究ではリスクアバース「累積報酬」や実施法に注目しています。「無線が途切れ、送りたい情報が届かないエラー」をリスクとし、無線独特の通信の誤りを超えるような観測状態、行動や報酬の設定になっています。さらに個々のIoTデバイスが、そういったリスク回避の行動を取りますが、最小限のコストで、多数のエージェントが協調することによって、より良いアクション、よりリスク回避する無線通信のアクションを選択するとい



リスクアバース強化学習を活用したマルチ無線アクセス制御技術の概要



った、そういった部分は我々のコアな技術アイデアです。

——工場での通信とデバイスの今後は？

河村 IoTデバイスでは様々な無線通信規格が使われます。例えば、Wi-Fiは、誰でも手軽に使える長所がある反面、干渉が多く、通信品質の安定性が低い場合があります。セルラー通信は干渉は少ないものの、工場内まで十分な電波が届くように基地局が整備されるとは限らないので、電波が来ない可能性もあります。また、周囲の一般ユーザも使用する周波数帯ですから、通信性能が担保できるとは限りません。ローカル5Gは自分達で基地局を設置でき品質も良好ですが、Wi-Fiに比べると「コストが高い」というデメリットがあります。現在のように、さまざまな電波や通信方式を使った多様な通信が共存し続ける状況は、今後も変わらないと思われます。

金子 6Gでさらに重要になってくるのがエネルギー利用効率で、それは通信性能とそれに必要な消費電力、コストのバランスを表す指標となっていて、本研究の一つの重要なターゲットとなっています。特にIoTではバッテリーが限られていますので、いかに小さな出力で、工場内での高度な無線通信やセンシングで高い品質を保証するかという点はエ

ネルギー効率にも関係しています。

各々の強みを共同研究に持ち寄り、さらなる高みへ

——今回のSICORPの研究プロジェクトでは、どのような協働になるのでしょうか？

金子 日本側アカデミアとしてはNII、基幹アルゴリズムの考案、評価を推し進め、企業としてはNTTにご参画いただき、省電力軽量エッジAI技術の創出の議論、適用アーキテクチャの検討、実験評価等を行います。また、フランス側ではアカデミアとして、IRISA/レンヌ大学が省電力AI圧縮技術開発、エッジAIソフトウェア・ハードウェア設計を担い、音響機器の企業であるWavely社が音響センサのIoT、環境変動抽出、異常検知のためのAI技術開発を進めます。

——それぞれに個性と持ち味のある4機関が、プロジェクトとして新しい何かを生み出すということですね。

河村 まさにSICORPのポリシー「イコールパートナーシップに基づく国際共同研究」を実現し、一つの国では出来ない国際共通課題を解決し、日本の科学技術力を強めることに貢献できるのではないかと考えています。

金子 私は無線システムの研究、それも干渉制御・性能解析・プロトコル設計という通信品質に大きく影響する部分の研究ですが、特に数理モデルによる研究を行ってきたので、どちらかといえば基礎研究です。NTTとは2018年から共同研究を続けていて、さまざまなユースケースを想定し、応用に向けた研究開発が可能になりました。その結果トップ国際

会議や論文誌発表のみならず12件の特許出願、うち4件は登録済みという、ある意味協創の成果を实らせました。

河村 NTTは企業ですから、どのように実装して、どう実用化してサービスに結びつけるかというところが主眼となります。金子先生がされているような要素技術の基礎研究は、我々はなかなか手が伸ばせないところです。それで、2018年から2021年にかけて「自律分散型無線



システムにおける機械学習を用いた環境推定技術」、2021年から2024年にかけて「マルチ無線アクセス環境における、高信頼用通信のための利用無線リソース最適制御技術」の2つのテーマで共同研究を続けてきましたが、それを踏まえての今回の国際プロジェクトです。

私たちも研究を通じて、ますます無線通信を発展させていきたいです。

金子 今はまだ、単純なIoTデバイスでも作動可能な省電力エッジAIを搭載した無線通信やセンシングシステムは、世の中に展開されていませんが、数年後には、工場の中に限らず、実現している社会を見ることができると期待しています。私も、SICORPでのNTTおよびフランスの2機関との協力で、その発展を目指していきます。

サプライチェーン リスクを 予測AIで 可視化

災害の多発や、紛争、パンデミックなどにより、サプライチェーンのリスクはより複雑化し、増大しつつある。そんな問題を起点として取り組むのが、東京海上ディーアール株式会社と国立情報学研究所(NII)による共同研究だ。同研究では、予測AIの技術を使い、取引先で起こったショックが対象企業にどう波及するかを定量的に把握。この成果が社会にどんな価値をもたらすのかを、研究を担う2人に聞いた。

佐藤 遼次

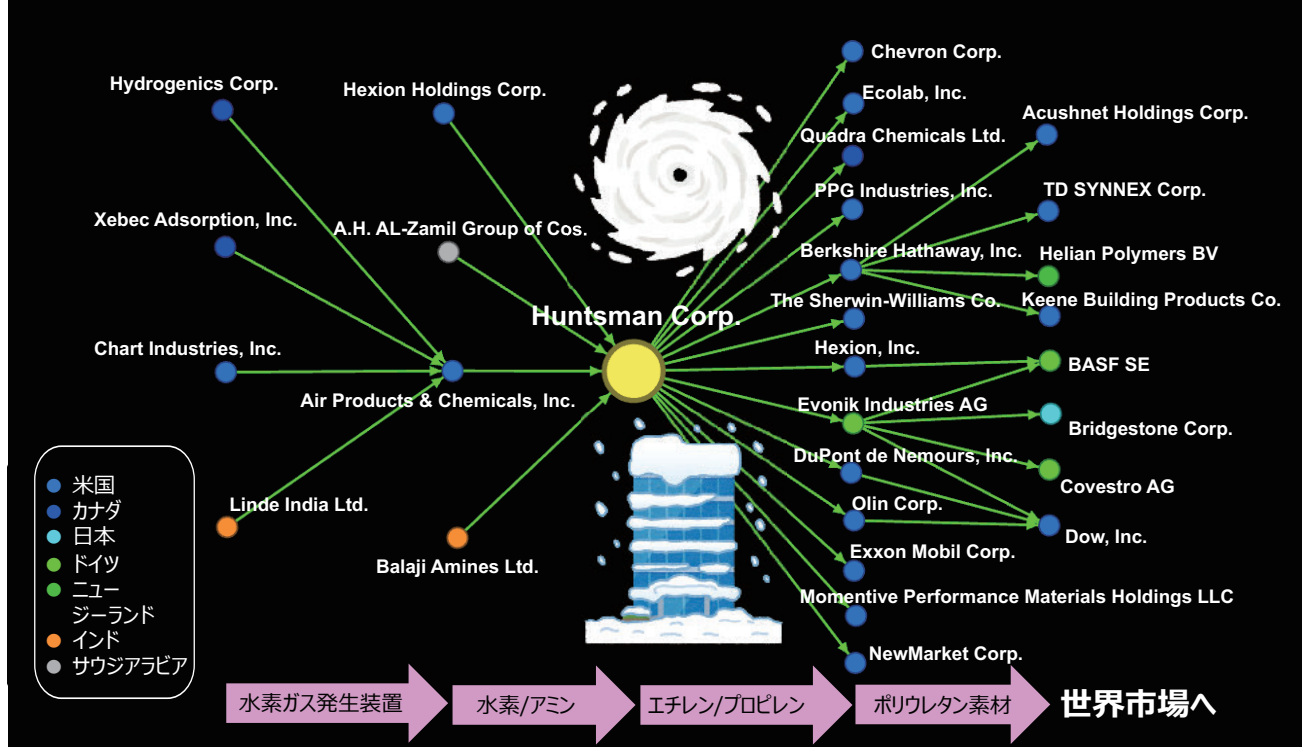
SATO, Ryoji

東京海上ディーアール株式会社
企業財産本部データビジネス創発ユニット 主任研究員

水野 貴之

MIZUNO, Takayuki

国立情報学研究所
情報社会相関研究系 准教授



FactSet Supply Chain Relationships 2021より推定

——共同研究を行うことになった経緯を教えてください。

水野 私がグローバルサプライチェーンの研究に従事したのが、2016年のことです。世界中で異常気象が頻繁に起きるようになり、かつ企業の取引ネットワークのグローバル化が進んだことで、グローバルチェーンのどこかで起こったショックが地域や業種を飛び越えて波及しやすくなった。それこそ2005年にアメリカで起こったハリケーン・カトリーナや2011年の東日本大震災などでも、調達物資の不足や輸送インフラの停止などの影響が、災害から遠く離れた場所の企業にも及びました。

そうした影響を見える化しようと2018年よりスタートしたのが、東京海上ディーアールとの共同研究です。私たちが携わる情報学では、インターネット上でどう情報が伝播するかを研究してきました。それをサプライチェーンに置きかえ、災害、地政学的リスク、パンデミックなどで、負の影響がどう伝播するのかを情報学の手法で追いかけてみようというのが今回の取り組みです。

佐藤 東京海上ディーアール企業財

産本部では、火災・自然災害リスクの定量評価およびコンサルティングを、顧客企業に提供しています。自然災害などに関する国や自治体等のデータを、顧客企業の情報と組み合わせでリスクを評価し、その結果を防災対策や適切な保険設計などのリスクマネジメントに活用いただくというものです。

一方、顧客企業のサプライチェーンとなると、複雑で秘匿性の高い情報でもあるだけに、精緻に把握することが困難です。そもそも顧客企業の方でも、直接取引を行う「一次取引先」までしか管理できていないケースも多い。したがって、災害などによるサプライチェーンの途絶といったリスクについては、大きく仮定を置くなどして評価せざるを得ません。その部分に大きな課題を感じていました。

そんな中で水野先生は、外部から取得できるビッグデータを使い、そこからAIでインサイトを引き出す研究をされていて、その技術を当社のもつ事故や災害の知見とかけあわせることで、これまでとはまったく違う角度から顧客企業のリスク評価を行えるのではないかと考え、共同研究を

ご提案しました。

取引先のショックのうち 約17%が対象企業に波及

——具体的に、どんな研究に取り組みましたか？

佐藤 取り組みの一つが、災害などによる負の影響が取引関係を介して波及する「ショックの波及」を、定量的に把握することです。ショックの波及が起きること自体は、先行研究で実証されていました。ただ、従来のアプローチでは、その大きさを定量的に評価することはできません。そこで今回、新たなアプローチとして構築したのが、AIの機械学習を用いた予測モデルでした。

水野 具体的には、このような手法を採りました。

まず、約38万件の取引関係を収録したグローバルのサプライチェーンのデータセットから、サプライヤー・カスタマーの双方を含めた対象企業の一次取引先を抽出する。続いて、抽出した対象企業および一次取引先に、「売上成長率」「国」「業種」などの情報を付与した上で、対象企業の売上成長率を、それ以外の付与した

情報から予測するモデルを構築する。そして、構築した予測モデルの挙動を分析し、取引先の売上成長率の低下が対象企業にどう影響をおよぼすかを調べる。なお、AI学習の手法は、「CatBoost（決定木の発展型アルゴリズム）」を採用しました。同手法で構築したAIモデルを、部分従属プロットという手法で「注目する情報(=取引先の売上成長率)だけを変化させたときに、予測対象値(=対象企業の売上成長率)がどう変化するか」を可視化した形です。

その結果、取引先の売上成長率が下がると対象企業の売上成長率も下がる傾向にあり、取引先のショックを100とすると、平均して約17%が対象企業に波及することが示されました。なお、取引先がサプライヤーかカスタマーか、あるいは製造業か非製造業かによる顕著な差はなく、ショックの波及の大きさはおおむね同水準でした。

佐藤 その後、この結果が実際のショックと大きく乖離していないか検証するために、2012年にアメリカ東海

岸を襲った「ハリケーン・サンディ」のケースと比較しました。結果、AIモデルで導いた一般的なショックの波及の大きさは、実災害におけるショックの波及と比べても大きな差異はなく、妥当な結果であると評価できました。

今回の予測モデルによって、売上成長率にもっとも影響するのは、その年や国、業種などに依存するマクロ経済的なトレンドであると同時に、取引先の業績といったミクロな情報も無視できないレベルで影響することが示されました。AIモデルを構築・分析することで、このようなマクロ要因と切り分けて取引先のショックの影響を定量的に把握できたことが、今回の共同研究の大きな意義だと感じています。

サプライチェーンの全容をAIで推定するシステム

——今回の研究成果を、企業や社会はどう享受できますか？

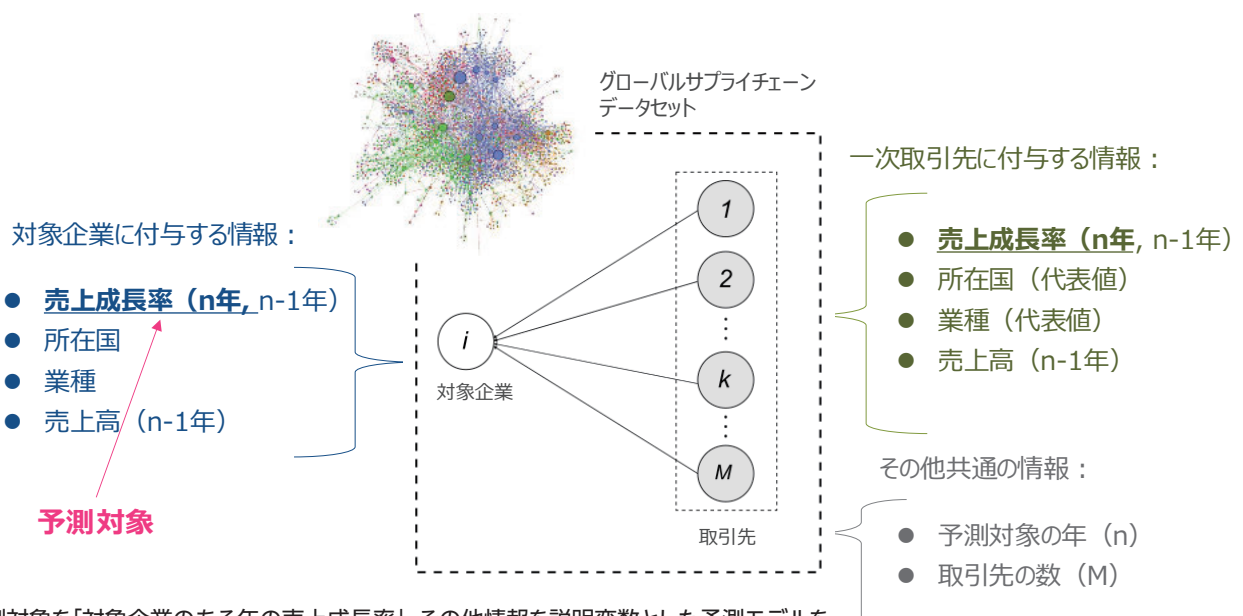
佐藤 社会実装の場の一つとなるのが、当社が手掛けるサプライチェーンコミュニケーションサービス

「Chainable(チェイナブル)」です。

Chainableは、企業のサプライチェーンの可視化やリスク管理を行えるプラットフォームで、2023年5月にローンチしました。顧客企業は、自社や取引先の登録拠点の付近で災害が発生したときにアラートを受け取れ、リスクのある拠点が可視化されます。また、自社や取引先の担当者に影響確認のタスクの一括発信やチャットなどの双方向コミュニケーション、ファイル共有などができ、有事での速やかな対応と平時からの対策や関係強化が行なえます。このChainableにおける、サプライチェーンのリスク評価の部分に、共同研究の成果の実装を検討しています。

とくに現代では、前述のようにサプライチェーンが複雑化し、当事者企業も一次取引先以外の経路を把握できていないことが少なくありません。そこで非常に有用になるのが、AIを使って膨大なサプライチェーンを推定する水野先生のシステムです。

——サプライチェーンを推定するシステムとは？





水野 まずは企業が開示している取引先情報を活用し、足りない部分に関しては、取引の経路をAIで予測するという技術です。これを使うことで、サプライチェーンのどこかで寸断が起こったときに、別の経路で代替できるのか、それとも代替が利かないので別の対応が必要になるのかといったことも予測できるようになります。

サプライチェーンの解析で難しいのが、加工や組立などを通してモノの形が変化すること、そして多くの企業が複数の製品を作っていることです。したがって、ネットワークで負のショックが起こったとしても、企業や製品によっては影響が出ないケースもある。リスクを過度に評価することは企業に不必要な対応を促すことにもなりかねないので、そこをきちんと区分けしてリスクを抽出する点に注力しました。

佐藤 たとえば有機化学製品を製造販売するテキサス州のグローバル企業、ハンツマン・コーポレーションは、およそ25カ国に50以上の製造拠点をもち、膨大なサプライチェーンを有します。同社は、2021年に米国メキシコ湾岸を襲った大寒波およびハリケーンでエチレン生産の大部分が停止したことにともない、機能製品部門で計画外のダウンタイムが発生しました。一方、先端材料部門、テキスタイル部門など他の部門での影響は限定的でした。膨大なサプライチェーンにおいて、ショックの影響範囲を適切にとらえるには、その製品を製造する際に上流と下流で何が行われ

るのかをたどれる水野先生の技術が重要なカギとなります。

水野 また、取引金額を推定することも非常に大切になってきます。サプライチェーンは、図で表すとただ線がつながっているだけですが、実際の線の太さ=取引金額は、各取引により大きく異なります。したがって、それも機械学習を使いながら予測していく。そのようにいろいろな観点から、有事の損失がいくらになるかの予測精度を高めることに、日々一緒に取り組んでいます。

取引を行う前に 取引先のリスクを可視化

——今回の研究を通して、どんな価値を提供していきたいですか？

佐藤 この技術により、関連取引だけを抽出してショックの影響を予測できるだけでなく、災害などから遠く離れた場所での影響も予測できます。たとえば、先ほどのハンツマンのような基礎化学品メーカーの製造が止まり、それによってポリウレタンなどの中間財や樹脂製品などの最終財の供給が不足する可能性があるといったリスクを、直接には取引をしていない企業にも提示できるというのは、とても有益ではないでしょうか。

また取引を行う前の段階で、取引先のリスクを推定することも、重要なテーマであるにとらえています。事前に把握することで、詳細な調査をかけるといった一歩踏み込んだアプローチも可能になります。

災害や地政学的リスクはもちろんですが、人権や環境対応、ガバナンスといった観点も、サプライチェーンのリスクをはかる上では不可欠です。そういった部分にもしっかり対応しな

がら、企業があるべき姿のサプライチェーンを築く助けになればなど。

共同研究で培った知見や技術を最大限活かすことで、多くの企業にサービスを使っていただける。それにより、多くの企業のサプライチェーンに関する知見が蓄積され、それを活用することで、より付加価値の高いリスク情報をご提供できる。ぜひ、そんなサイクルを回していきたいです。



水野 特定の目的のもとで取引先を選定すると、当然、特定の地域に固まりがちです。それによって生産効率は高まるかもしれませんが、一方で有事のリスクも高まってしまう。そのリスクを事前に可視化することで、ちょっと分散させようとか、寸断が起こったときのために自社で供給のバッファを作っておこうといった事前の対応や備えが行えます。

それと個人的には、AIでサプライチェーンの行く末を予測するといった領域に、大きな興味関心をもっています。この先、どんなリスクが、どのくらいの確度で起こるのか。それに対して、たとえばこの政策を打つことで、サプライチェーンのリスクはどう改善されるのか。そんなふうな、近未来の日本や世界の展望を見すえながら最適な政策を見つける研究を、将来的には進めていきたいです。

より自然で高強度な 「音声匿名化」を 求めて

個人情報の保護・秘匿が、以前よりもますます厳しく問われている現在。テレビ放送で取り上げられるさまざまな「生の声／証言」にも、一層の注意が払われるようになった。研究者と放送局が連携し、新たな音声匿名化技術の実装を果たした事例について日本放送協会（NHK）のご担当者、および国立情報学研究所（NII）の山岸順一教授に聞く。

高石 真美子

TAKAISHI, Mamiko

日本放送協会 デザインセンター
音響デザイン部

山田 正幸

YAMADA, Masayuki

日本放送協会 デザインセンター
音響デザイン部

山岸 順一

YAMAGISHI, Junichi

国立情報学研究所
コンテンツ科学研究系 教授

音声の匿名化が抱えていた問題

特定の個人を識別できる情報にはさまざまなものがある。氏名や住所などの属性情報から、顔つき、指紋、声などの生体情報など——これらはその人の「存在」を証明するのに重要である一方で、悪用されればその人を傷つけることにもなるため、大切に守らなければならない。

テレビ放送においても、主に報道番組やドキュメンタリー番組などでは、一般の方々がコメント、証言などに登場するケースは多々あるが、その際に、さまざまな事情から、匿名性を担保したい場合もある。特に声に関しては、従来は、高く／低く変調して元のをわかりにくくするという手法がよく採られていた。しかし、そこには問題もあった。

「例えば、映像的にはモザイク等を掛ける。音声は変調した素材を重ねて複雑にし、聞きづらはテロップ表記で補う。従来はそれしかないということで、当たり前のようにその表現を使ってきたのですが、勇気をもって発言して頂いているにもかかわらず、時にはネガティブな印象を与えてしまったり、あるいは人工的過ぎて人間性への配慮が不足してはいないか。本当にこの手法しかない

のか。いろいろな現場でディレクターや編集担当者など制作スタッフと話し合うこともあり、そんなモヤモヤを常に抱えて来ました」

福祉関係のドキュメンタリーなどを多く担当してきたという、NHKデザインセンター音響デザイン部の高石真美子氏は、このように語る。

そもそも、単純な変調の場合は逆の操作を行えば元の音声が発元出来てしまう。NHKの場合は、多重にピッチ変更を行ったうえでエフェクトを加えるなどの加工も行っていたが、それでも非可逆性は万全とは言えず、前述のように加工後の質に不満もあった。一方ではフランス国立音響音楽研究所(IRCAM)で開発されたシステムの試用も行われていたが、これらにも、変換に時間がかかったり、あるいは扱う技術者によって変換結果に違いが出たりと、なお実用性に問題が残った。

そんななか、NHKの開発部局主催の勉強会に出席していたNIIの山岸順一教授に会ったことをきっかけに、相談を持ち掛けたことが、新たな音声匿名化の実装へと繋がっていったという。

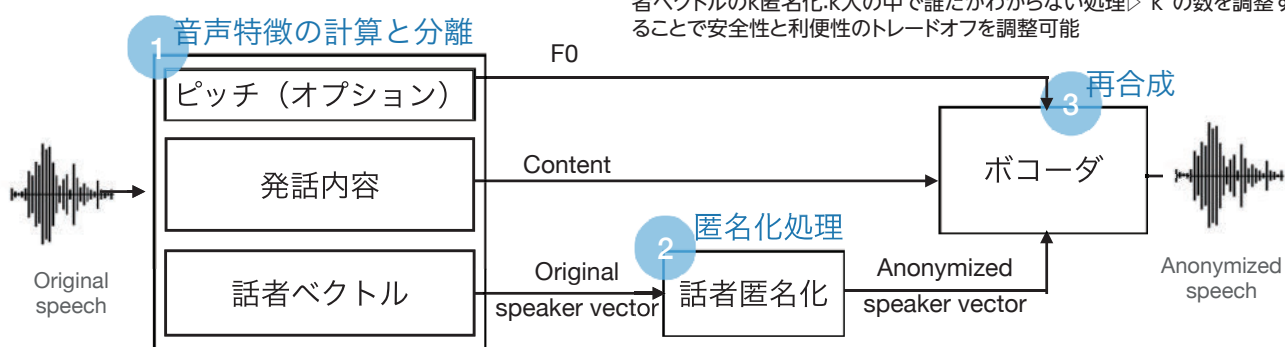
システム開発と番組制作への導入

山岸教授は、20年以上にわたり

音声合成の研究に携わってきた。特にボイス・クローンあるいはデジタル・クローンと呼ばれる、特定の個人の声を再現する技術などを手掛けてきたという。

「近年になってディープ・ラーニングの技術の発展などもあり、ボイス・クローニングに関してはだいぶ高いレベルで実現できるようになってきました。そうした時に、新たに浮上ってきて興味を持ったテーマが2つあります。1つは、ボイス・クローニングによって高度に再現された声、いわゆるディープ・フェイク音声を悪用された場合に、それをどのように見破ればよいのかということ。もう1つは、誰の声でも合成できるのであれば、逆にそれを私たちのプライバシー保護にも使えるのではないかと、ということです。特に後者に関しては、個人情報保護法の改正などもあり、音声もパーソナルデータとして、第三者への配布時には匿名化しなければならないといったことも義務付けられるようになっていました。しかし一方で、では具体的に、どのような処理を行えば匿名化が果たされるかということは、明確化はされていませんでした。もちろん、ノイズ化してしまえば匿名化できるでしょうが、内容も失われてしまう。利便

話者自動匿名化プログラム「VocalGuard」の仕組み



①音声信号を次の2つもしくは3つの特徴系列に分離(a)発話内容(b)ほとんどの個性を含むと考えられる話者ベクトル(c)必要ならば、基本周波数(ピッチ)②話者ベクトルのみを加工し③音声を再合成する加工方法は話者ベクトルのk匿名化:k人の中で誰かわからない処理で「k」の数を調整することで安全性と利便性のトレードオフを調整可能



性とプライバシー保護の両立を図らなければならない。そうしたことも整理しつつ、技術開発を進めていたのです」(山岸)

NHKからのアプローチは、まさにそのようなタイミングでのことだった。最初の意見交換が2021年8月のこと。その後実際に、山岸教授の技術を元に音声匿名化システムの実装に取り掛かった。

「そして、このシステムをいよいよ初めて番組で実用したのが昨年(2023年)末のことです。以来、存在が口コミで広がって、実際に使ってみたいという要望も含め、問い合わせが集まってきています。やはりこうした課題に対するニーズは非常に高かったのだということを、あらためて実感しています」(高石)

実装にあたり、手法として選ばれたのが「k匿名化」と呼ばれる技術だった。

具体的には、ある発言(音声)を、発話内容と抑揚、個人性の3つに分解する。このうち、個人性のみを匿名化し再構成するのだが、この際に、音声データのプールを活用し、個人が特定化される確率をk分の1以下に変換する。この際、kの値は設定次第で変更が可能で、これを大きくすれば、その分、匿名性は向上する。「NHKの方からお話を頂く前から研究開発を進めていたのですが、当初、2019年頃に作ったシステムは英語対応のものでした。実際にNHKで使って頂けるものにするには日本語対応にする必要がありまし

た。単純に日本人の音声データに入れ替えればいいというものではなく、『k人の日本人のなかで、匿名性が保たれるようにする』点での調整も必要でした」(山岸)

より使い勝手のよい匿名化を

現時点で実装している匿名化システムでは、男女別やkの値によって、何段階かの変換モードをプリセットし、利便性を高めている。「完全に別の声に置き換わってしまうのは避けたい」という要望にも応えるため、kのなかに発言者本人の声も含める選択枝も用意されている。「やはり、科学的に高い不可逆性が担保されているということ、そして自然でありながら、ほどよい加工感があり『本人ではない別の誰かに似てしまう』懸念が少ないことも、実際に使ってみて感じる良い点ですね。現場でも好評で、新たな選択枝として定着していく可能性を強く感じています。今後は、東京だけでなく全国の拠点でも広く利用できるようになればと思っています」(高石)

「実は、匿名化を図ることとは違うのですが、ドラマ方面からもとても魅力的な技術で、たとえば近未来が舞台の作品で役者さんの声にこの加工を用いることによって、AIやデジタル世界の声を表現したり、別人格に生まれ変わらせたりなど、演出手



法の一つとして使わせてもらおうかという話もあります」(山田)

一方、実用化されていく中から見えてきた課題もある。すでに山岸教授のなかで、今後の改良のポイントとして設定しているものが2つあるという。

「一つは、より安全性／匿名性を向上していくことです。まずは音声のプールを増やし、より大きなkの値を設定し使用できるように、ということを考えています。これに関しては特に技術的な障壁は無く、新たなデータ集積の手間の問題だけなので、早々に手を打ちたいと思っています。もう一つは、実際に使ってみてNHK側から指摘されていることですが、おそらくインタビューを録っている状況によって、処理後の音声に少し外国語風イントネーションが掛かっているように聞こえる場合があるという問題です。これについては、音声を発話内容と抑揚、個人性の各要素に分解する際の切り分け方の工夫で対処できるのでは、と感じています。やはり、単なる研究で、論文を書くことだけを思うなら必要なくても、実際に使用する場面で重要なことがいろいろ出てきます。それは私にとっても大いに参考になります」(山岸)

「声の個人情報」の秘匿に関して、新たな一歩を開いたこの技術。今後のさらなる社会実装への展開を期待したい。

アカデミアと インダストリーの ストーリー

最前線の研究に携わるアカデミアの研究者と、産業界で製品を創り出すエンジニア。いずれも世の中を変える技術を生み、育てていく達人だが、両者の間に見えない隔壁があるという。その壁とはいったい何か。アカデミアとインダストリーで、手探りながら協創関係を構築してきた二人に、現在に至るストーリーを聞く。

安里 彰 + 五島 正裕
富士通株式会社 国立情報学研究所



安里 彰

ASATO, Akira

富士通株式会社

先端技術開発本部システム第一開発部

Industry

■「まずは話してみよう」から始まる
関係

——まずはお二人の簡単なお経歴、
そして交流のきっかけについて伺え
ればと思います。

五島 私は東京大学の准教授から、
ちょうど10年前に教授として国立情
報学研究所に移ってきました。研究
テーマは学生の時から一貫してコン
ピュータ・アーキテクチャを手掛けて
おり、その絡みで、安里さんともお付
き合いが始まったという感じです。
安里さんは富士通の中でもちょっと
珍しい方で、学会にも非常にコミット
して頂いていて、それがそもそもの
きっかけです。

安里 私は一貫して富士通の所属
です。3年程前に定年を迎えていま
すが、現在も嘱託という形で関わっ
ています。もう少し詳しく言うと、前
半は富士通研究所という研究子会
社に所属し、研究といっても大学に
おけるそれとは異なり、富士通の製
品開発を視野に入れた試作のような
ことをやってみました。その後、2007
年からはスーパーコンピュータ「京」
のプロジェクトが始まったのを契機
に組織変更があり、私を含むグルー
プが事業部門に移籍、そのまま定年
までという流れです。

特に前半の、研究寄りの仕事をし

ていた時期に、上司の勧めもあって、
学会に関わり、研究会の幹事のような
役も務めるようになりました。私の
性格もあるのかもしれませんが、
会社の中にずっといるよりも、そうし
て外の先生方と交流を持つのが楽し
くて、事業部門に移ってから、そんな
関係が続きました。事業部門にい
ながらも、積極的に外に出ていき、
アカデミア方面の方々ともお付き合
いがあり、顔も覚えて頂いていたと
いう点では、ちょっと珍しい社員だっ
たかもしれません。

もっとも、五島先生ご本人に関し
て言えば、学会関係の仕事をするな
かでお顔は存じていましたが、当初
はその程度でしたね。

——そうした状況から、どのように
協力関係を進めて行ったのでしょうか。

五島 もともとは、東京大学のある
研究室の懇親会で、若手の先生と、
その研究室出身の、安里さんの部
下の方との間で話が進み、何か共同
研究のようなものは出来ないかとい
うことになった。その時に、私にも声
が掛かって、その集まりに加わった
というのが、直接のきっかけだったと
記憶しています。

安里 もっともその時は、はっきりと
「共同研究」という形にまではなっ

ていませんでしたね。共同研究という
と、企業側からきちんとお金を出し
てアドバイスを貰ったり、ということ
になると思うのですが、その時の集
まりはそこまでではなかった。組織と
してきちんと契約を結び、ということ
ではなく、ほぼ個人レベルで集まっ
て、「何か面白いことある?」というよ
うな感じでした。

五島 そもそも当時の状況でいう
と、富士通側にも大学とそこまで緊
密に、何か一緒にやるという雰囲気
はありませんでした。確かに社内では
独自に研究開発は進めているのだ
ろうけれど、外から何かを取り入れ
ようとしている風にも見えなかった、
というのがあります。

安里 確かにそれはあったと思いま
す。私としては、富士通の中にいな
がら大学の先生方とも付き合う立場
にいて、いま五島先生が仰ったよう
なことも感じていました。社内を見
渡して、「この人たち、もっと外を見
て、先端の研究成果を取り入れてい
けばいいのに」なんて思うこともあ
りましたし。

そんな中で、「こんな集まりを持と
うと思うがどうか」という話が回っ
てきた。これはいい話が出てきた、渡
りに船だ、と。できる範囲で積極的
に進めて行きたい、これで社内の皆

にももっと刺激を与えたいと思ったのです。

五島 実を言うと、私は当初、あまり期待はしていなかったんです。一緒にいった先生の中には、企業で作っているプロセッサの特性の実測値が欲しいとか、そんな希望を持っている方もいましたが、私は、どうせ何も出てこないだろうと。

実際に、企業が行っている研究開発の中身は喋れないことが多くて、共同研究レベルになるときちんと秘密保持契約(NDA)も結んで、「どこまで話せる／話せない」をきっちり決めるのですけれど、その時はそこまでも行っていません。

ただ、当初の想像に反して意外なほどに歓迎してもらって、話も聞いてもらって——。最初のうちは、半年に一度くらいでしょうか。会合を持って、今、アカデミアではこういうことを考えている、研究室ではこんな研究をしている、という話を聞いてもらうということが多かった。

安里 そうですね。ですから、「ギブ&テイク」という観点でいえば、その頃は我々はほとんど貰うばかりで、研究の最先端の様子を聞かせていただくという感じでした。

五島 それでも、それ以前に富士通

に対して抱いていたイメージからすると、信じられないくらいオープンだったんですよ(笑)。

それが、私がNIIに移ってきてすぐの頃だったと思います。そうしたことを重ねてきて、2020年頃くらいよいよ本格的な共同研究の話に移っていくことになります。

■交流から得られるもの

——当初の「共同研究」まで至らないミーティングの中で、具体的な研究成果のやり取りなどは無くても得られるものというのは何なのでしょう。

五島 アカデミアの立場からすると、公知の範囲であっても、自分自身のやっていることを誰かに話すこと自体に意味があるというか、楽しい。それも、実際にプロセッサを作っている人たちに聞いてもらえることに価値がある。もちろん、アカデミアといっても個人差もあると思いますが、特に我々のように、インダストリーに近い立場にいる者としては、やはり自分の考えていることが世の中に出て行って欲しいという希望があるわけです。ただ実際には、それを世の中に繋げていくパスがなかなか存在しない。そんな時に、本格的な

共同研究ではなく、雑談レベルであったとしても、自分の喋る内容が何らかの影響を与えて、それで富士通のプロセッサが少しでも良くなっていくのであれば、それはハッピーなことだと思うのです。

安里 私自身は、インダストリーとアカデミア、その両者が見える位置にずっといて、「見えているのに、接点がない」ということに、ストレスを感じていました。特に開発部門のエンジニアたちには、アカデミアの世界で今何が語られているかに、もっと興味を持ってほしい。もちろん本当は、自分で社外に積極的に出て行って交流を広げてほしい。

そこで得た情報をどう使うか、判断はもちろんエンジニア自身がやればいいのだけれど、とにかく、井の中の蛙になって、「今何が起きているのか」を知らずに開発を進めているのでは、特にプロセッサのような分野では、絶対に世界に後れを取ってしまうと思っていました。そこで、こうした機会を突破口にして、さらによりよい関係を築いて行ければと思っていましたね。

例えば海外の場合は、企業が積極的に著名大学に寄付講座などを作り、そこで優秀な学生を集めて、



五島 正裕

GOSHIMA, Masahiro

国立情報学研究所

アーキテクチャ科学研究系 教授



Story

どんどん共同研究を進めて論文を出し、賞を取ったりしている。当然ながら、そうした成果は製品にも反映されていく。最近は富士通もそういった動きが少し出てきてますが、当時はまだまだで、いずれはうちもそんな風にしていきたい、という思いがありました。

五島 もともとアカデミアとインダストリーの交流があまり盛んではなかったことについては、日本企業の「自前主義」のような部分も背景にあるのではと思っています。これは程度の差はあれ、日本企業の多くが持っている気がしているのですが、技術を外から買ってくることはせず、全部自社内で産み出すんだという気持ちが非常に強い。

しかしそれでは、これからの時代は戦っていけません。少なくともコンピュータ・アーキテクチャの分野では、アメリカと日本では研究コミュニティの規模が開いてしまっています。研究者の層の厚み、関係する人材の数自体が、おそらく一桁以上違います。ですから、研究を行うにしても、よほどすごい戦略を立てて当たらないと、一角に食い込んでいくことすら難しい。そのためには、やはりアカデミアとインダストリーの連携は欠かせないと思います。

■見えてきた課題と今後の展望

——協力関係を進めて行くにあたり、特に難しさを感じる点、今後の課題等について伺えますか。

五島 当初に比べると、現在はだいぶ状況は改善されてきていると思うのですが、それでも秘密主義的な部分は相変わらず残っている印

象です。

実際に今は共同研究も行っているのですが、例えば研究部門と結んでいるNDAと、事業部門と結んでいるNDAとは異なっていて、研究者のレベルでは、事業部がやっていることをどこまで話していいのかが、手探りだったりします。そしてそれを判断するのに、ものすごく時間がかかったりする。アメリカのような、大学研究室と企業との一体感のようなものにはだいぶ遠いと感じます。

安里 ただ、アメリカの場合は、日本のように組織によって人間が厳然と分かれているのではなく、頻繁に人材面でも行き来があるとか、同じ人が両方に所属しているか、そういった環境もあるのではと思います。おそらく、会社とか大学とか、そういった枠組みを越えたところで研究者・技術者のコミュニティが出来上がっていて、その中で活発に情報のやりとりも行われている。そうしたカルチャーは日本には無いものですね。

もうひとつ、私の側から言うところ——。私自身は最近では共同研究等に関わっていないので具体例とかではないですが、やはり企業側の意識の問題は残っていると思います。ここまでで、徐々に改善したものはあると思いますが、やはり企業の閉鎖性のようなものはある。たぶん、現場レベルでは問題意識を持っている人はいると思うのですが、やはり最終的には経営判断としてgoを出す人たちがもっと変わってほしい。

五島 いや、特にここ5年くらいでしょうか。だいぶ変わってきたと思いま

すよ。それは安里さんがきっかけを作ってくれたのも大きいと思う。やはりミーティング等で、設計部隊が現実困っている問題についてアドバイスして、それが功を奏してうまく行った、というようなことが少しずつ積み上がっていった、それにつれて信頼関係も増している気はしますね。

いずれにしても、そのような、アカデミア側からの貢献部分、いわば“グッド・プラクティス”を、規模的にも数的にも、もっと積んでいく必要はあると思っています。本当のところを言えば、アカデミアは企業のことなど考えずとも、ひたすらに良い研究をして、それを企業側が見出して応用してくれるというのが理想的だと思いますが、なかなかそうはいかない。やはり、アカデミア側から、「こうなるとよくなりますよ」ということを、積極的に発信していくことも必要だと思っています。実際には、ここ最近は富士通側から「ここ、どうしたらいいでしょう」と聞いてくることも増えて、昔はそんなことはなかったもので、隔世の感を覚えているところではあります。

もっとも、10年かけてここまで変わってきた……ではあと10年でどこまで変わるのか。10年経ったら私も定年ですから、このペースで大丈夫なのかというのはありますね(笑)。もっとドラスティックに変わって欲しいところではあります。

安里 私自身はもう定年を迎えていて、囑託での退職もあと数年というところなので、自分がどうこうというのはありませんが……。それでも、これまで10年かけて積み上げてきたものが、よりうまく回っていくようになってほしいですね。

NII ニュース・トピックス一覧

NEWS TOPICS

期間: 2024/3/1 (金) から
2024/7/31 (水) まで



各ニュースの詳細は
オンラインでご覧になれます。

www.nii.ac.jp/news/2024

Facebook

[www.facebook.com/
jouhouken/](https://www.facebook.com/jouhouken/)

X (旧Twitter)

twitter.com/jouhouken

YouTube (音が出ます)

[www.youtube.com/user/
jyouhougaku](https://www.youtube.com/user/jyouhougaku)

情報犬ビットくん

X (旧Twitter)

twitter.com/NII_Bit

NII Today

ご意見募集中

www.nii.ac.jp/today/iken

メルマガ

購読随時受付中

www.nii.ac.jp/mail/form

NII Todayはオンラインでバック
ナンバーもご覧になれます。

www.nii.ac.jp/today

冊子版は、学校図書室、研究機
関など団体への配本を受け付け
ています。お申込み、配本停止は、
kouhou@nii.ac.jp 広報へメー
ルでご連絡ください。個人への
送付はおこなっておりませんが、
冊子は東京千代田区一ツ橋の
学術総合センタービル1階に常
設されており、ご自由にお持ち
になれます。

ニュースリリース

NEWS RELEASE

2024

- 7/18 コンピュータサイエンスパーク2024 7/31(水)開催!〜プログラミング思考を
学べる遊び場がいっぱい 5年ぶりの名物企画「研究100連発」も!〜
- 7/17 学術情報ネットワーク(SINET6)のデータ流量を日本地図にマッピングしたプロ
ジェクション作品を展示ー「動いてあたりまえ」のネットワークインフラの重要性を
可視化ー
- 6/13 組合せ遷移分野における15年来の未解決問題を肯定的に解決ー理論計算機科
学分野のトップカンファレンス「STOC 2024」にてNIIの平原 秀一 准教授らの
論文採択ー
- 5/30 生成AIでソクラテスに人生相談できる!? 西洋古典学の飛躍的發展に寄与す
る対話システムを開発
- 5/23 NIIの活動を幅広く紹介!「NIIウィークス2024」ーオープンハウス(6/7)、学術情
報基盤オープンフォーラム(6/11-13)、ジャパン・オープンサイエンス・サミット
(6/17-21)を連続開催ー
- 5/22 OERリポジトリ試行版の提供開始ー大学が開発した教材の横断検索が可能にー
- 4/30 大規模言語モデル「LLM-jp-13B v2.0」を構築ーNII主宰LLM勉強会(LLM-jp)
が「LLM-jp-13B」の後続モデルとその構築に使用した全リソースを公開ー
- 4/18 NIIと東工大が日本語版の大規模言語モデルの研究開発における連携協定を締
結ー生成AIモデルの透明性・信頼性の確保と社会実装の加速化への取り組みー
- 4/10 ISO 34502の自動運転車危険シナリオを数学的に定式化ー安全性保証タスク
の自動化・効率化により自動運転の社会受容を促進ー
- 4/1 国立情報学研究所に「大規模言語モデル研究開発センター」新設ー国産LLMを
構築し、生成AIモデルの透明性・信頼性を確保する研究開発を加速ー
- 4/1 NIIに「トラスト・デジタルID基盤研究開発センター」新設
- 3/28 学術検索基盤の研究開発の最先端に触れるーCiNii Labsサイトの公開ー
- 3/27 NIIとOpenAIREが協力協定を締結ー協働で研究基盤の研究開発を進め オ
ープンサイエンスの推進に貢献ー
- 3/18 JSTとANRの戦略的国際共同研究プログラム(SICORP)「日本ーフランス国際
産学連携共同研究」(エッジAI)研究課題への採択ーAIを用いた自律的な無線
アクセス制御技術の研究開発の加速ー
- 3/4 GakuNin RDM データ解析機能に新機能を追加ーMATLABでより高度な数
値解析が可能にー
- 3/1 PtM:合成音声付き動画教材作成システムー実証実験をスタートー

受賞

AWARDS

2024

- 7/3 堀込 泰三さん(総合研究大学院大学、水野研究室)、水野 貴之 准教授(情報社
会相関研究系)らの論文が人工知能学会2023年度研究会優秀賞を受賞
- 6/10 Alex Spies日本学術振興会外国人特別研究員、井上 克巳 教授(情報学プリン
シプル研究系)らによる論文がTechnical AI Safety 2024で最優秀ポスター
賞を受賞
- 6/3 佐藤 一郎 教授(情報社会相関研究系)が電子情報技術産業協会会長賞を受賞
- 5/28 大規模言語モデル勉強会(LLM-jp)がAAMT長尾賞を受賞
- 4/22 来るべき情報技術の社会的信頼を担う数理的ソフトウェア研究で文部科学大臣
表彰・科学技術賞(研究部門)を受賞。NIIの蓮尾 一郎 教授、石川 冬樹 准教授、
京都産業大学の勝股 審也 教授が共同受賞
- 4/9 「メタな視点に基づく平均時計算量の研究」で NIIの平原 秀一 准教授が若手科
学者賞を受賞ー令和6年度の文部科学大臣表彰ー
- 3/15 佐藤 諒平さん(総合研究大学院大学・高須研究室)がDEIM 第16回データ工学
と情報マネジメントに関するフォーラムにて学生プレゼンテーション賞を受賞
- 3/5 鯉淵 道紘 教授(アーキテクチャ科学研究系)に電子情報通信学会がフェロー・称
号を授与
- 3/5 平原 秀一 准教授(情報学プリンシプル研究系)が第11回ヤマト科学賞を受賞

Essay

日本の科学技術の明るい未来のために

片岡 洋

KATAOKA, Hiroshi

国立情報学研究所 所長代行/副所長

今年2月、日本のGDPがドイツに抜かれて世界4位に転落したと報じられた。日本経済の「失われた30年」とよく言われるように、1人当たりGDPは、2000年に日本は世界2位だったが、2010年に18位、2022年に33位まで転落している。

科学技術分野の指標でも懸念すべき状況が示されている。文部科学省科学技術・学術政策研究所(NISTEP)の「科学技術指標2024」によれば、日本の論文数(2020～22年平均)は、中、米、印、独に次ぐ5位だが、注目度が高い論文(他の論文に引用された回数が各分野で上位10%に入るトップ10%論文)の数で、前回と同じ過去最低の13位(前回、イランに抜かれて12位から転落)。日本は20年前は4位だったが、それ以降順位が下がり続けている。また、引用数が極めて多い「トップ1%論文数」も、前回と同じ12位(前回、10位から転落)。

もちろんこれらの指標が全てではないが、1995年の科学技術基本法(2021年科学技術・イノベーション基本法に改正)制定後、5年毎の科学技術基本計画(今は科学技術・イノベーション基本計画)策定、2001年の総合科学技術会議(今は総合科学技術・イノベーション会議)設置等、科学技術振興に力を入れてきたはずなのに、一

体どうしたことだろうか。

研究開発費総額や研究者数は、米中の近年の大幅な増加に対して、日本は伸びが小さく停滞しているが、順位は世界3位を維持している。とすると、それらが研究力低下の要因ではないように思える。

令和4年版科学技術・イノベーション白書は、日本の大学を対象としたNISTEPの分析を紹介し、近年の論文数の減少要因として、教員の研究時間割合の低下、教員数の伸び悩み、博士課程在籍者数や直接的に研究実施に関わる費用の停滞といった要因を挙げている。また、研究パフォーマンスを高める上での制約に関するアンケートを基に、具体的な制約として、教育負担や大学運営業務によって研究時間が確保できないことや、基盤的経費の不足等によって研究資金が確保できないことを挙げている。

日本の科学技術関係予算は、2024年度に当初予算4.9兆円となり、2001年度(3.5兆円)に比べ4割増えている(これに補正予算等を加えた額は近年特に大幅に増え、2022、23年度は9.5兆円近くになっているが、補正予算では研究者を雇用することができない)。特に競争的資金は数も予算も増えたが、多様な研究シーズを生み出す元になる基盤的経費は減り、不足して

いる。人材問題も深刻で、博士課程の入学者数は2003年度をピークに長期的に減少傾向にある(2023年度は4.4%増に転じた)。大学の本務教員に占める40歳未満の割合も減少し続けている。

研究力低下については、既に様々に議論され、施策も講じられてきているが、まだ具体的な成果が上がっているようには見えにくい。私も科学技術行政に携わり、持ち場では最善を尽くしてきたつもりだが、局所最適が必ずしも全体最適にはならないということも大きいのではないか。今こそ国全体として叡智を結集して総合的・中長期的な検討を徹底的に行い、全体最適を実現して日本の科学技術の明るい未来を切り拓く努力が必要であろう。一方、NIIに関して言えば、科学技術において情報、データ、AIの役割は今後ますます大きくなり、NIIはその中核にいるので、果たすべき役割は大きいと思う。例えば、生成AIについては、NIIでもいち早くLLM-jpという活動を立ち上げて国産LLMの開発等オープンな共創の取組を進め、今や1,600人(2024年7月現在)を超す産学の人々に参加頂いている。本年4月からは文部科学省からの予算により大規模言語モデル研究開発センターを設置したところである。このような取組をますます発展させていくことが重要である。

情報から知を紡ぎ出す。

NII

国立情報学研究所ニュース:NII Today 第103号 令和6(2024)年9月
発行:大学共同利用機関法人 情報・システム研究機構 国立情報学研究所
〒101-8430 東京都千代田区一ツ橋2丁目1番2号 学術総合センター

本誌についてのお問い合わせ:総務部企画課 広報チーム
EMAIL:kouhou@nii.ac.jp

発行人:黒橋 禎夫

編集委員長:越前 功 監修:武田 浩一 アドバイザー:山本 浩幾
編集委員:池畑 諭、金子 めぐみ、込山 悠介、水野 貴之、竹房 あつ子(ABC順)
編集/AD:岸本 治恵 制作:創言社、日精ピーアール
表紙イラスト:市村 譲

