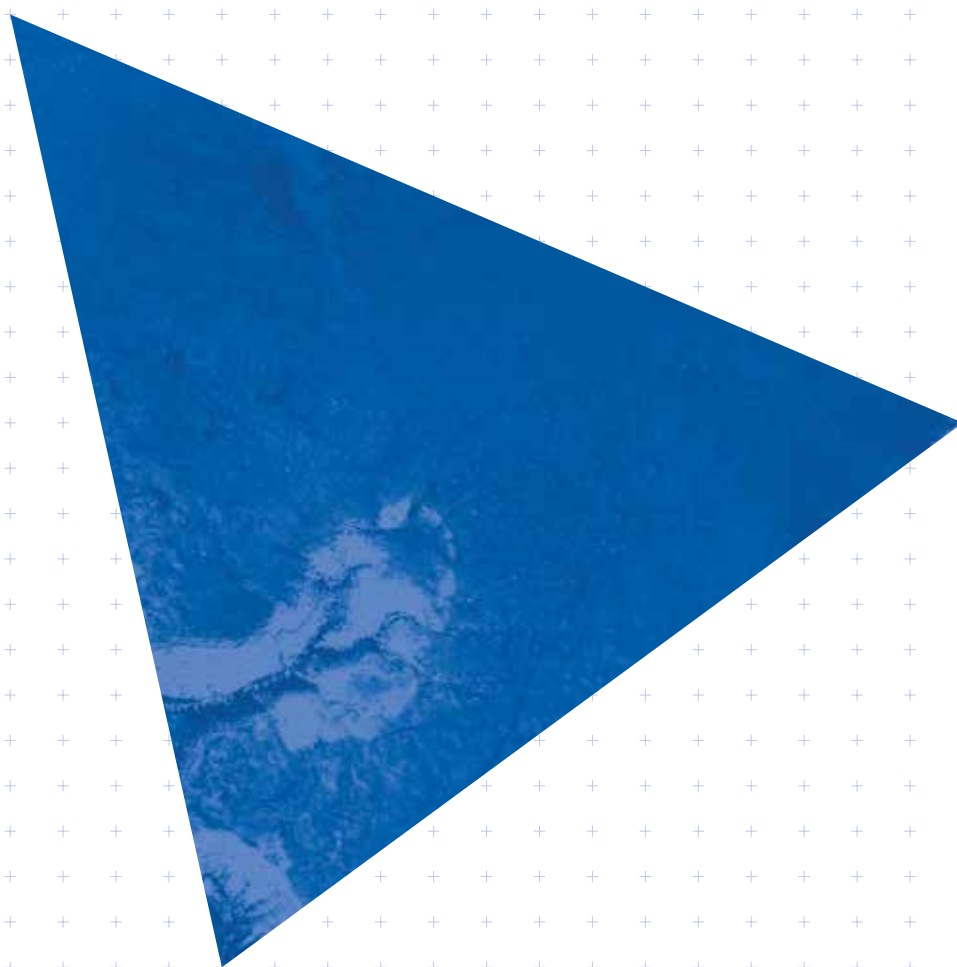


俯瞰ワークショップ報告書

エージェント技術

2022 年 1 月 26 日開催



エグゼクティブサマリー

本報告書は、2022年1月26日に開催した「エージェント技術俯瞰ワークショップ」の内容をまとめたものである。

「エージェント」とは、自ら判断し行動する主体をコンピューターシステムとして実現したものである。実社会のさまざまな活動・現象をマルチエージェントシステムとしてモデル化・シミュレーションする「マルチエージェント」の研究開発、人間と音声・言語で対話する「対話エージェント」の研究開発、人間の心理・認知面も踏まえ、対話だけでなく外観・振る舞い等も含めた総合的な「ヒューマンエージェントインタラクション (Human-Agent Interaction : HAI)」の研究開発などが、活発に取り組まれているトピックである。このような研究開発によって、情報サービスやロボットとの新しいインターフェースの創出や、さまざまな社会活動・社会現象の分析・予測による社会課題解決につながると期待されている。

本俯瞰ワークショップでは、上記の3つのトピック「マルチエージェント」「対話エージェント」「HAI」を中心に、外部有識者9名を招き、注目すべき研究動向や今後の方向性・重要課題について議論した。その中では、ビッグデータと機械学習（深層学習、強化学習、逆強化学習など）の活用による実世界行動モデリングやエージェント設計論の進化、マルチエージェント間や人間－エージェント間のインタラクションにおけるさまざまな要因の取り込みが進んでいることが確認された。また、社会・人間の活動の理解や社会課題解決への貢献に向けて、「マルチエージェント」「対話エージェント」「HAI」の間の連携・融合強化の必要性や、マルチエージェント社会シミュレーションや大規模学習モデルなど計算リソース面の課題も論じられた。

科学技術振興機構（JST）研究開発戦略センター（CRDS）は、科学技術に求められる社会的・経済的なニーズを踏まえて、国として重点的に推進すべき研究領域や課題、その推進方策に関する提言を行っている。そのための基礎資料として、隔年で「研究開発の俯瞰報告書」を公開している。2021年版（2021年3月公開）がその最新版であり、次期版は2023年3月公開を予定している。この俯瞰報告書で取り上げる対象分野の一つとして、人工知能・ビッグデータ分野があり、その中に含まれる研究開発領域の一つが「エージェント技術」領域である。

今回開催した俯瞰ワークショップは、2023年版の作成に向けて、その「エージェント技術」領域に関して、最新動向を把握・アップデートすることを主な目的とした。冒頭、CRDSからイントロダクションを兼ねて、2021年版の俯瞰報告書に「エージェント技術」の概要・トピックを紹介し、その後、「マルチエージェント」「HAI」「対話エージェント」のそれぞれについて、外部有識者から話題提供をいただいて議論を行った。最新動向や今後の方向性の把握に加えて、各取り組みの間の知見共有・連携にも良い機会となった。

目次

1		エージェント技術動向：俯瞰報告書 2021 年版から	1
2		マルチエージェント領域（1）	8
		発表 2-1：伊藤 孝行	8
		発表 2-2：荒井 幸代	13
		討議	17
3		マルチエージェント領域（2）	19
		発表 3-1：野田 五十樹	19
		発表 3-2：森永 聡	22
		討議	26
4		H A I（Human-Agent Interaction）領域	28
		発表 4-1：山田 誠二	28
		発表 4-2：小野 哲雄	31
		討議	34
5		対話エージェント領域	36
		発表 5-1：東中 竜一郎	36
		発表 5-2：駒谷 和範	40
		討議	43
6		総合討議	45
付録			49
		付録 1 ワークショップ開催概要	49
		付録 2 参考文献・参考情報	50

本ワークショップに参加いただいた外部有識者

氏名	所属等
伊藤 孝行	京都大学大学院情報学研究科 教授
荒井 幸代	千葉大学大学院工学研究院 教授
野田 五十樹	北海道大学大学院情報科学研究院 教授
森永 聡	NECデータサイエンス研究所 上席主席研究員
山田 誠二	国立情報学研究所 教授
小野 哲雄	北海道大学大学院情報科学研究院 教授
東中 竜一郎	名古屋大学大学院情報学研究科 教授
駒谷 和範	大阪大学産業科学研究所 教授
大沢 英一	公立はこだて未来大学大学院システム情報科学研究科 教授

1 エージェント技術動向：俯瞰報告書 2021年版から

科学技術振興機構（JST）研究開発戦略センター（CRDS）は、科学技術に求められる社会的・経済的なニーズを踏まえて、国として重点的に推進すべき研究領域や課題、その推進方策に関する提言を行っている。そのための基礎資料として、隔年で「研究開発の俯瞰報告書」を公開している。2021年版（2021年3月公開）^[1]がその最新版であり、次期版は2023年3月公開を予定している。

この俯瞰報告書で取り上げる対象分野の一つとして、人工知能・ビッグデータ分野があり、その中に含まれる研究開発領域の一つが「エージェント技術」領域である。今回開催した俯瞰ワークショップは、2023年版の作成に向けて、その「エージェント技術」領域に関して、最新動向を把握・アップデートすることを主な目的とした。「エージェント技術」の研究開発に取り組んでいる有識者9名を招き、注目すべき研究開発動向や今後の方向性・重要課題について、話題提供をいただいて議論を行った。

本章では、有識者からの話題提供（第2章以降）に先立ち、俯瞰報告書の2021年版に記載した「エージェント技術」の概要を振り返る。

人工知能・ビッグデータ分野におけるエージェント技術の位置付け

図1-1は、台頭した技術やエポックを時系列にプロットした、人工知能・ビッグデータ分野の時系列俯瞰図である。3回の人工知能（AI）ブームを通して、理論の革新、応用の革新、社会との関係に取り組みが広がり発展してきたことを示している。

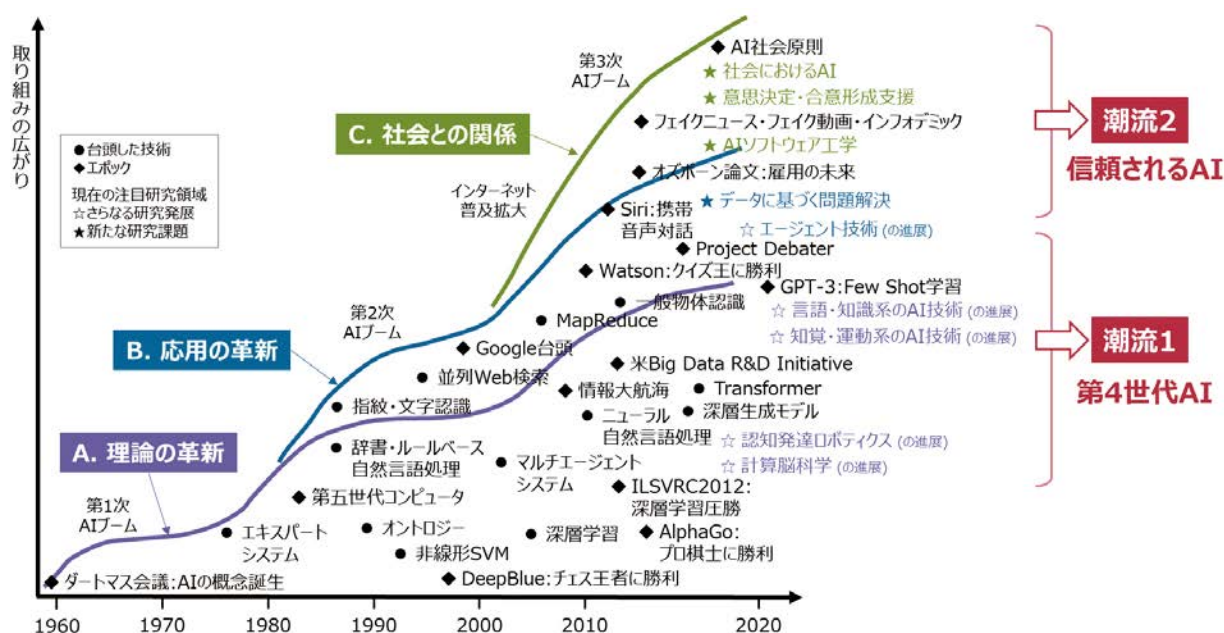


図1-1 人工知能・ビッグデータ分野の時系列俯瞰図^{[1][2]}

このような発展を経て、現在、「第4世代AI」と「信頼されるAI」という2つが、研究開発の潮流となっている。「第4世代AI」の潮流は、AI技術そのものの精度・性能の向上に取り組む研究開発で、深層学習を中心とした機械学習ベースの「第3世代AI」の次の世代を目指す取り組みである。一方、AIが社会に広がったことで、社会受容の面からAIシステムに対して、精度・性能とは異なる、安全性・信頼性・説明性などの要請が強まった。それに応えるための研究開発が「信頼されるAI」の潮流である。

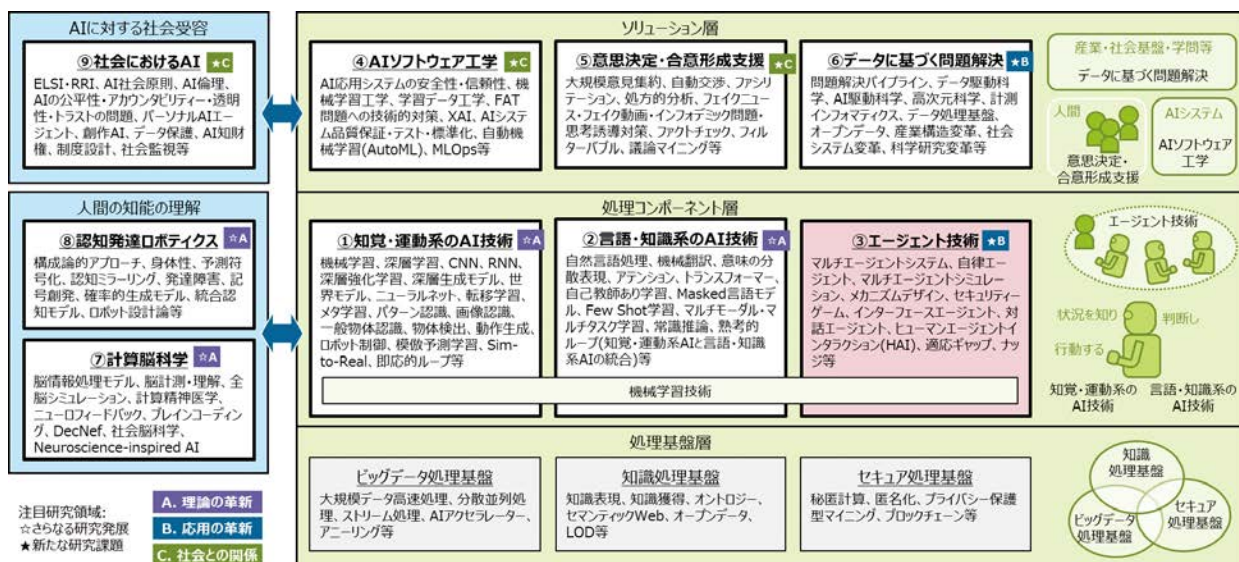


図1-2 人工知能・ビッグデータ分野の構造俯瞰図^{[1][2]}

そこで、2021年版の俯瞰報告書では、この2つの潮流を踏まえて、注目すべき研究開発領域を設定した。図1-2に、人工知能・ビッグデータ分野の構造俯瞰図を示す。主に処理コンポーネント層が「第4世代AI」、ソリューション層が「信頼されるAI」に関わる研究開発領域となる。

処理コンポーネント層には、注目する研究開発領域を3つ設けたが、それらの関係は図1-2の右端に図示したように、1つのAIの内部処理を「状況を知り、判断し、行動する」と分解したとき、「状況を知り」「行動する」に対応するのが「知覚・運動系のAI技術」領域で、「判断し」に対応するのが「言語・知識系のAI技術」領域である。一方、AIと外部とのインタラクション（AIとAI、AIと人間）に関わるのが「エージェント技術」領域である。

ソリューション層にも、注目すべき研究開発領域を3つ設けた。それらの関係についても図1-2の右端に示したが、処理コンポーネント層の技術群を用いたAIシステムの作り方に関するものが「AIソフトウェア工学」領域、AIシステムと人間との関係に関わるものが「意思決定・合意形成支援」領域、AIシステムによる産業・社会基盤・学問などの変革に関わるものが「データに基づく問題解決」領域である。

また、処理コンポーネント層を作る上で、人間の知能の理解が参考になるが、そのような研究開発として、「計算脳科学」領域や「認知発達ロボティクス」領域がある。ソリューション層に対しても、AIに対する社会受容に関する研究開発として「社会におけるAI」領域がある。

このような人工知能・ビッグデータ分野の全体観のもと、今回のワークショップでは「エージェント技術」領域にフォーカスする。なお、上述の9つの研究開発領域の動向や研究開発の2つの潮流については、別報告書^{[1][2]}で詳しく説明している。

エージェントの定義・特性

まず、エージェントを「自ら判断し行動する主体をコンピューターシステムとして実現したもの」と定義した。これは広い意味では、AI そのものだとも言える。また、エージェントが持つ特性として、自律性・自発性（自分自身の動作の目標を設定して動作）、社会性（他のエージェントと協力して組織を構成して問題解決を実行）、反射性（種々の変化や変動を察知して適応的に動作）などが言われている。

エージェントに関する研究開発で主に扱われているものとして、(a) 自律的メカニズム（自律エージェント）、(b) 複数主体の協調（マルチエージェントシステム）、(c) 社会活動・現象のシミュレーション（マルチエージェントシミュレーション）、(d) 人間とのインタラクション（インターフェースエージェント）などが挙げられる。

以降、(a) (b) (c) を「マルチエージェント」としてまとめ、一方、(d) の「インターフェースエージェント」を、自然言語などによる対話をベースとした「対話エージェント」と、外見や応答の仕方も含めた総合的なインタラクションを扱う「HAI（Human-Agent Interaction）」に分けて説明する。

マルチエージェント領域の研究開発動向

マルチエージェントによって、さまざまな社会活動・現象のモデル化とそれを用いたシミュレーション・予測を行うことができる。また、それを通して、社会活動・現象のより正確な理解や意思決定への活用が可能になる。



図1-3 社会活動・現象のマイクロ・マクロループ

これに関連して言及しておきたいのは、社会活動・現象のマイクロ・マクロループである。つまり、図1-3に示すように、各個人の行動の集積が社会全体の動きとなる一方で、社会全体の動きを受けて各個人の行動が変化する。したがって、ある期間のビッグデータから機械学習を用いて規則性を見いだしたとしても、将来の動きはその規則性から外れていってしまう。これに対して、マルチエージェントシミュレーションを使えば、構成論的アプローチによって、ダイナミクスを捉え、多面的な理解・予測が可能になる。

表1-1に、俯瞰報告書2021年版に記載した本領域の研究開発動向を要約した。2021年版では、注目動向としてメカニズムデザインとセキュリティゲームを取り上げた。

表1-1 マルチエージェント領域の研究開発動向(2021年版の要約)^{[1][2]}

領域の意義	●さまざまな社会活動・現象のモデル化とそれを用いたシミュレーション・予測 ●それを通して、社会活動・現象のより正確な理解や意思決定への活用
これまでの研究開発動向	●1980年代：分散人工知能と呼ばれ、分散協調問題解決の方法論が中心 ●1990年代：各問題解決器が自律性を備えたマルチエージェントシステムに（BDI（信念－願望－意図）モデル、強化学習、サブサンプションアーキテクチャー等） ●2000年代：経済パラダイムとの好相性から、ゲーム理論・メカニズムデザインのオークションやマッチング問題への適用で実用成果 ●2010年代：社会実装が強く志向され、電子商取引、電力マネジメント、交通シミュレーション、ワイヤレスセンサーネットワーク等へ適用 ●社会活動・現象を扱うマルチエージェントシステムの中に人間も含めてモデル化（人間に与えるインセンティブの設計、人間の限定合理性の考慮等） ●トップランク国際会議としてAAMAS、コンペティションAANC、国内ではJAWS等
注目動向	●メカニズムデザイン：ゲームのルールを設計することで、利己的なエージェントの行動を間接的に制御 ・オークションやマッチングで実績 ・自動メカニズムデザイン、オンラインメカニズムデザイン等の発展拡張 ●セキュリティーゲーム：攻撃者から重要施設を守るための警備員の適切配置問題 ・警備側・攻撃側ともに人員リソースは有限で、互いに相手の動きを読み合いながら効果的に人員配置 ・ロサンゼルス空港警備計画、野生動物保護レンジャー配置計画、公衆衛生問題等で実績
今後の方向性	●マルチエージェント社会シミュレーション：さまざまな社会的問題解決への適用拡大 ●シミュレーションに用いるエージェントモデルの課題・限界：人間の判断・行動のモデル化の困難性、現実世界の多様性・不確実性 ●現実世界の問題への活用の方角・課題：正確な予測よりも問題・リスクの可能性と対策検討への活用、そのための条件を変えた反復による計算量増大への対応

HAI領域の研究開発動向

HAIは、人間の心理・認知にも目を向け、自然言語だけでなくエージェントの外見・非言語情報まで含めた総合的なインタラクションに着目する。その際、仮想エージェントとのインタラクションだけでなく、物理的な実体を持つロボットとのインタラクションや、エージェントを介した他の人間とのインタラクション（Computer-mediated Communication）も研究対象とする。そして、HAIは、外見や応答の仕方も含めて、ユーザーそれぞれが心理的・認知的に受け入れやすいインターフェースエージェントの設計方法を与えてくれる。これは、インターフェースエージェントの自然性・親密性・効率性などの向上に結びつく。

表1-2に、俯瞰報告書2021年版に記載した本領域の研究開発動向を要約した。HAIは日本発の比較的新しい研究領域である。2021年版では、注目動向としてHAI設計論の進展を挙げた。

対話エージェント領域の研究開発動向

対話エージェントによれば、普段使っている言葉（自然言語）でコンピューターに指示や意図を伝えられる。また、ハンズフリー・アイズフリーで使える。タスク指向型対話で実用化が先行し、現在では、非タスク指向型対話（雑談）中心のチャットボットも実用化されている。

表1-3に、俯瞰報告書2021年版に記載した本領域の研究開発動向を要約した。自然言語処理分野と同様に、深層学習を用いた大規模言語モデルを用いた対話システムが開発されるようになり、2021年版では、それらの注目システムとして、Google Meena、Facebook Blender、OpenAI GPT-3を挙げた。

表1-2 HAI 領域の研究開発動向（2021年版の要約）^{[1][2]}

領域の意義	●外見や応答の仕方も含めて、ユーザーそれぞれが心理的・認知的に受け入れやすいインターフェースエージェントの設計 ●自然性・親密性・効率性等の向上
これまでの研究開発動向	●人間の心理・認知にも目を向け、自然言語だけでなくエージェントの外見・非言語情報まで総合的なインタラクションに着目 ●日本発の比較的新しい研究領域 ・国内でHAIシンポジウム（2006～） ・日本発の国際会議HAI（2016～） ●さまざまな利用場面が広がる3種類のHAI研究対象（人間－仮想エージェント、人間－ロボット、人間－エージェント－人間）、学際的研究が進展
注目動向	●HAI設計論：人間とインタラクションするエージェントの設計法は未確立だが、考慮すべき事項や指針として以下等 - エージェントのデザイン：外見、表出表現、機能レベル等の適切な選択 ・適応ギャップ：擬人化における期待と実際の機能・能力の差異リスク ・メディアの等式：メディア世界と現実を無意識のうちに同一視 ・ナッジ：自発的な行動を促す刺激、悪用を避けて倫理面に配慮 ・意図スタンス：物理／設計スタンスに加えて意図スタンスも理解して設計
今後の方向性	●HAI設計論の確立：受け止め方の個人差を吸収する設計方法論、ケーススタディー研究から統一的な方法論へ基礎的な知見も積み上げ ●キラーアプリケーションの探索やHAIが効果的な応用実証：医療、介護、福祉、教育、ビデオゲーム等 ●実用化の進んでいるマルチエージェントシミュレーションや対話エージェントにHAI研究の知見展開

表1-3 対話エージェント領域の研究開発動向（2021年版の要約）^{[1][2]}

領域の意義	●普段使っている言葉（自然言語）でコンピューターに指示や意図を伝えられる ●ハンズフリー・アイズフリーで使える
これまでの研究開発動向	●まずタスク・利用シーンを限定した形でタスク指向型対話が実用化 ・コンタクトセンターでの問い合わせ応答、旅行会話の音声翻訳等（Nuance、IBM、NEC、NICT等） ●スマートフォンやスマートスピーカーが音声対話アシスタントを搭載 ・2011年 Apple Siri、2012年 NTT DoCoMoしゃべってコンシェル、2014年 Microsoft Cortana ・2014年 Amazon Echo/Alexa、2016年 Google Home/Assistant、2017年 LINE Wave/Clova等 ●非タスク指向型対話（雑談）中心のチャットボットは新たな顧客接点にも ●パイプライン構成から深層学習 End-to-End 構成へ ・大規模データ学習による対話エージェント ●トップランク国際会議：SIGDIAL、ICMI 等、コンペティション：Alexa Prize 等
注目動向	●大規模データ学習による対話エージェント：Transformer型の大規模言語モデルが対話システムにも次々と（2020年） ・Google Meena：会話341GBからパラメーター26億個を学習したチャットボット ・Facebook Blender：会話15億件からパラメーター94億個を学習したチャットボット ・OpenAI GPT-3：テキスト3千億件からパラメーター1750億個を学習、対話文に限らず少数例示からさまざまな文章を生成可能
今後の方向性	●音声対話における技術課題への対応：文脈の一貫性、応答の多様性、外部知識の参照、キャラクターの一貫性、制御性の確保、評価方法・評価尺度 ●視覚情報等も含めたマルチモーダル情報統合・実世界連動：（例）ロボットに対して「それを赤い箱に片づけて」「缶コーヒーを買ってきて」と指示する

表1-4 エージェント技術の国際比較（2021年）^{[1][2]}

国・地域	フェーズ	現状	トレンド	各国の状況、評価の際に参考にした根拠など
日本	基礎研究	○	→	HAIは日本発の研究分野である。自動交渉エージェントに関するANACでも日本の貢献が見られる。対話分野の国際会議では一定数の論文が採択されているが、エージェント基礎研究全般として研究者層や国際学会での存在感はまだ弱い。
	応用研究・開発	○	↗	オンデマンド交通・避難計画等の交通・人流シミュレーションでの実用化が見られる。エージェント研究者も関わっている人狼ゲームやロボカップ等のコンペティションも日本が先導している。
米国	基礎研究	◎	→	他のAI分野と同様に研究者層が厚く、エージェント分野のAAMASや対話分野のSIGDIALやICMIの採択論文数も圧倒的1位である。
	応用研究・開発	◎	→	セキュリティゲームでの実用化をはじめ、エージェント技術の先端的な応用が進んでいる。大規模学習による対話エージェントでは圧倒的優位性を示しており、Facebook、Amazon、Microsoft等からの論文発表も多い。
欧州	基礎研究	◎	→	米国に次いで研究者層が厚い。特に英国の存在感がある（AAMASの採択論文数では米国に次いで英国が2位をキープしている）。欧州はマルチモーダル対話も比較的強い。
	応用研究・開発	○	→	基礎研究に比べると、企業から目立った研究成果の発表は見られない。
中国	基礎研究	○	↗	AAMASの採択論文数は徐々に増加している。自然言語処理・対話関連はACL・EMNLPで米国と並んで二強。
	応用研究・開発	○	↗	AI分野全体では米中二強の状況であるが、対話分野（タスク指向・非タスク指向）で企業（Tencent、Alibaba等）の取り組み・発表が活発であることを除いて、エージェント分野全般での中国の存在感はまだそれほどではない。しかし、中国のAI応用は急速に拡大しており、エージェント応用研究面も伸びつつある。
韓国	基礎研究	△	→	AAMASの採択件数もごくわずかである。SIGDIAL・ICMIでも目立った活動は見られない。
	応用研究・開発	△	→	目立った活動は見られない。
イスラエル	基礎研究	◎	→	バルイラン大学等はメカニズムデザインやゲーム理論等の基礎研究でトップレベルである。
	応用研究・開発	△	→	企業から目立った研究成果の発表は見られない。

（註1）フェーズ

基礎研究：大学・国研などでの基礎研究の範囲

応用研究・開発：技術開発（プロトタイプの開発含む）の範囲

（註2）現状 ※日本の現状を基準にした評価ではなく、CRDSの調査・見解による評価

◎：特に顕著な活動・成果が見えている ○：顕著な活動・成果が見えている

△：顕著な活動・成果が見えていない ×：特筆すべき活動・成果が見えていない

（註3）トレンド ※ここ1～2年の研究開発水準の変化

↗：上昇傾向 →：現状維持 ↘：下降傾向

政策面の課題

俯瞰報告書 2021 年版では、政策面の課題として、次の2点を挙げた。

1 点目は、大規模コンピューティング基盤・データ基盤の整備・強化である。対話エージェントの大規模言語モデル学習やマルチエージェント社会シミュレーションなど、計算量が増大している。理化学研究所の「富岳」や産業技術総合研究所の「ABCI」のような共同利用計算環境の継続的な整備・強化が求められる。

2 点目は、産学連携および分野横断の研究開発推進である。実社会・実用現場でのデータ取得や検証のために産学連携の活用が有効である。また、HAI や対話研究では、さまざまな分野（AI、ヒューマンインターフェース、ロボティクス、認知科学、心理学、哲学など）の知見が必要で、学際的な研究コミュニティ作りが求められる。さらに、今後は、マルチエージェント研究、HAI 研究、対話エージェント研究の間の連携が重要になっていくので、その促進策も考えていく必要がある。

表1-4には、2021 年版時点のエージェント技術全体としての国際比較を示した。日本が先導・先行する取り組みもいくつかあるが、総合的にはトップポジションにあるわけではない。3つの領域ごとに挙げた今後の方向性への取り組みを推進するとともに、2つの領域間のシナジーを生み出していくことも重要であろう。

2 | マルチエージェント領域（1）

発表2-1：伊藤 孝行

マルチエージェントシステムとは

私が思っているマルチエージェントシステムの定義は、1つの研究分野であって、社会の特徴的な側面を切り出した計算モデルのデザインや評価を行う分野であるというものだ。人工知能（AI）が「一人の人間」の知能の計算モデルであるならば、マルチエージェントシステムは「社会」の知能の計算モデルだと考えている。社会の特徴的な側面としては、例えば、インタラクション、リソース配分、交渉、提携、合意、説得といったものがあって、計算モデルとして捉えられれば、それらの知的な自動化が可能になる。

その上で定義を述べると、「マルチエージェントシステムとは、問題を表現・解決するためのアプローチとして、複数の知的な主体（エージェント）が、共通または分散された問題を解くためのシステム、または、そのシミュレーターである」と考えている。

重要な特徴は、社会を対象としているので、個々の計算が非集中型である。さらに、個々のエージェントは個人的な合理性を持っている。個人的な嗜好があって、プライバシーを守ろうとする。つまり、情報を出せと言ってもなかなか出してくれない場合もある。というような面も考えて、言ってみれば自分勝手なエージェントたちをどうやってハーネス（Harness）していくか、というところに本質的な面白さがあると思う。

社会実装の進展と研究開発の3つの注目動向

以前は、マルチエージェント研究は結構理論的なもの、例えばモデルだけとかタイプロblemが多かった。しかし、インターネットやスマートフォンの普及であるとか、計算機パワーの発展とか、深層学習に代表されるような知能処理の新しいアルゴリズムやモデルが出てきているといったことから、モデルを直接社会に実装できるフェーズに入ってきたと思う。社会実装の事前検証や社会の政策決定のために、マルチエージェントシミュレーションの重要性も増している。

研究開発の注目動向としては、①マルチエージェントシステムの合意形成のための交渉と協調^{[1][2]}、②マルチエージェント深層強化学習、③マルチエージェントシミュレーション、という3点を挙げたい。本日私は特に①の交渉や協調のトピックにフォーカスをして話したい。このあと、②は荒井先生、③は野田先生から紹介されると思う。

マルチエージェントの合意形成のための交渉と協調

マルチエージェントの交渉と協調に関するトピックを表2-1に示した。交渉と協調に関する研究が脈々と続けられているが、ゲーム理論がそれらの重要な数学的基礎になっている。それに基づいて、理論的に発展させる方向と、その発想を持って応用に進んでいく方向へと進んでいる。

基礎理論では、真実申告最良メカニズムは重要なトピックだと思うし、均衡解に対するMediated Equilibriumという新しいアイデアにも注目している。応用面では、自動交渉のほか、不完全情報ゲームのポーカーで人間に勝ったことも話題になった。セキュリティゲームも実世界に応用されていて、非常に重要だ。

表2-1 マルチエージェントの交渉と協調に関するトピック

基礎理論	ゲーム理論、メカニズムデザイン、提携ゲーム、シャープレイ値、ヘドニックゲーム、マッチング、投票、社会選択理論、ソーシャルネットワークなど ・ 真実申告最良メカニズム（Vickrey オークション、VCG など）：正直が最良の策になるオークションやメカニズム ・ Mediated Equilibrium ・ 深層学習による自動メカニズムデザイン、限界貢献ネットワークによる提携形成計算の効率化 ・ ソーシャルネットワークの Influence Maximization
理論をベースにした応用	自動交渉（バーゲニング／協力ゲーム） ・ 応用 競技会、実世界応用、合意形成、交渉シミュレーション
	不完全情報ゲームで勝つアルゴリズム ・ 応用 後悔最小化アルゴリズムによる相関均衡のポーカーへの応用
	セキュリティーゲーム（スタッセルブルグ均衡の応用） ・ 応用 空港警備、湾岸警備、鉄道警備、密猟警備などへの応用

自動メカニズムデザイン

メカニズムデザインは、マルチエージェントシステムの分野でも中心的なトピックの一つになっている。真実申告最良のオークションメカニズム（Incentive Compatibility）、つまり、正直に申告すると一番得する（損をしない）というメカニズムがコアになっている。第二価格オークション（Vickrey オークション）がよく知られている。真実申告最良の考えで設計していくとメカニズム自体が安定するので、Googleの広告のメカニズムなどに応用されていることはよく知られている。

注目している方向性の一つは、2002年にConitzerとSandholmから提案された自動メカニズムデザインである^[3]。真実申告最良は一つの制約であって、その制約を満たすようにメカニズムをデザインするのを自動化しようとアイデアは20年前から言われていた。その実現方法として、David Parkesのグループが2019年に深層学習を用いたRegretNetという手法を開発した^[4]。彼らはこれをオークションの収入最大化で成功させた。

このオークションの収入最大化は、1981年にMyersonが単一財の場合についてデザイン^[5]して以来30-40年間、解けていなかった。2人2財というケースでも解けなかった。実は私もこの問題に少し関わり、相互依存価値というのを導入することでこれにアプローチし、非対称情報を使うと収入が多少増えることは示すことができた。そして、2006年にAAMASというマルチエージェントシステムのトップカンファレンスでBest Paper Awardをいただいた^[6]。

RegretNetでは、真実申告最良の制約を満たすというのを、Regret値の最小化問題として扱い、深層学習を用いて解いた。財の配分を決めるネットワークと、支払額を決めるネットワークという2つを深層学習モデルとして訓練することで、データに基づいてメカニズムが自動設計される。これによって40年間、理論的には解けていなかった複数財の収入最大化が可能になった。ここで、Regret値を少しだけ緩めると収益が少し増えることも分かってきた。

Mediated Equilibrium

もう一つ、Mediated Equilibrium^[7]にも注目している。図2-1 左側の「囚人のジレンマ」はとてもよく知られている問題だが、それぞれプレイヤーが合理的で、かつ、コミュニケーションができない状態だと、両プレイヤーとも裏切るという均衡に落ちてしまう。ここで、Mediatorという選択肢を追加するとどうなるか（図2-1 右側）。Mediatorはどういうものかという、両プレイヤーがMediatorに合意するならば「協力する」という手を取るけれど、片方のプレイヤーだけがMediatorに合意した場合は「裏切る」という手を取るものである。そうすると、両プレイヤーともMediatorに合意するという均衡に至る。コミュニケーションがない状態でもプレイヤーはMediatorにという選択肢があると協力する可能性があることが示された。これはゲーム理論の本質的なところに踏み込んでいるように思える。

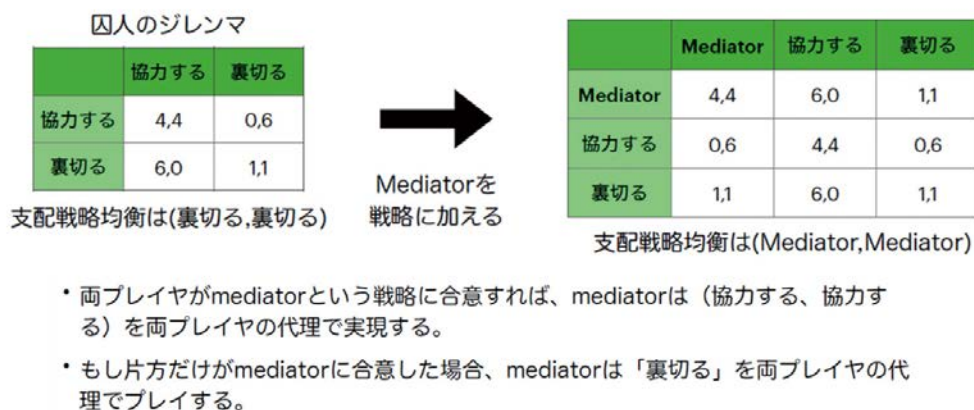


図2-1 Mediated Equilibrium

不完全情報ゲーム

チェス、将棋、碁は完全情報ゲームと呼ばれ、ゲーム木のどこにいるかが分かる。ご存知のように、AlphaGoが深層学習とモンテカルロ木探索を使って世界トッププロに勝利して話題になった。

一方、ポーカーは不完全情報ゲームであり、ゲーム木において、自分が今どのノードにいるかが一つに決められない。非常に状態数が増えて、計算量が増大し、完全情報ゲームに比べて、はるかに複雑になる。このようなポーカーでも、ゲームAIが世界チャンピオンに勝利した^{[8][9]}。ここでは、Regret Minimizationという方法がCounterfactual Regret Minimization（CFR）に拡張され、さらに、モンテカルロCFRという方法が開発された^{[10][11][12]}。

Regret Minimizationでは、先ほど説明したRegretNetと同じく、あるプレイヤーがある行動を選択したけれど、そのほかの行動を選択していたら、これくらい効用を得られたはずだという「後悔」（Regret）を数値化して用いている。このRegret値を最小化するようにゲーム木を更新していくと、相関均衡（緩いナッシュ均衡）に収束することが証明されている。

ポーカーのような不完全情報ゲームではプレイヤーがゲーム木のどのノードにいるかが分からないので、それを情報セットと呼び、情報セットまでの到達確率を統計的に計算して、各情報セットのCounterfactual Valueを求めた。そして、その平均を取ることでCounterfactual Regretというものを算出し、これを最小化するようにした。さらに、モンテカルロCFRは、モンテカルロ木探索にCounterfactual Regretを使う。モンテカルロ木探索では、ある分岐点からゲームの最終局面までシミュレーションを行うが、そのときにCounterfactual Regretを使って分岐確率を調整するようにした。このモンテカルロCFRでもゲーム木が非常に大きいので、オリジナルのゲームからゲームをより抽象化する制約緩和のようなことも考えられて、これ

に深層学習を使って自動化する Deep Counterfactual Regret Minimization というアイデアも出てきている。

このような技術によって6人制ポーカーでAIが勝利した。不完全情報ゲームについて紹介したが、ゲーム木の中で自分がどこのノードにいるか分からない状態があるというのは、社会における問題、さまざまな交渉場面でよく出てくることなので、いろいろなところに応用できる可能性がある。

また、ポーカーAIの発展は、毎年トップカンファレンスで行われたポーカーの国際競技会の貢献がかなり重要で、多くの先端的な研究者による独創的なアイデアが取り入れられることで上記が実現された。

セキュリティゲーム

もう一つ注目すべきトピックとしてセキュリティゲーム^[13]がある。これはゲーム理論におけるスタックルベルグ均衡を使っている。ゲーム理論における設定は非常に簡単で、スタックルベルグゲームには2人のプレイヤーがいて、片方のプレイヤーが戦略を取って、それを見た後にもう一つのプレイヤーが戦略を取るというものである。

空港の警備を例としてスタックルベルグゲームを説明する（図2-2）。警備対象としてTerminal 1とTerminal 2があり、警備を担当するポリス（Defender）はTerminal 1の方を重要視して、有限の人員を配置する。すると、攻撃者（Adversary）はそれを見て、Terminal 2に有限の人員を回して攻撃するというようなことが起きる。Terminal 1を一日の60%、Terminal 2を一日の40%守るといったような混合戦略も表現できる。さらに、Terminal 1の方が重要だと考える攻撃者もいれば、Terminal 2の方が重要だと考える攻撃者もいるという状況もあり得る。そのように複数タイプのプレイヤーがいる場合、タイプ分けしていくと、ベイジアンスタックルベルグ均衡の問題になる。そして、さまざまな条件下でのスタックルベルグ均衡解を求めるアルゴリズムが、AAMASなどのトップカンファレンスで多数提案されている。

- 空港の警備を考える。空港には2つのTerminal、1つのpoliceユニットがあるとする
- Terminal 1の方が、Terminal 2より重要とする。

		Adversary(attacker)	
		Terminal 1	Terminal 2
Defender (police)	Terminal 1	5, -3	-1, 1
	Terminal 2	-5, 5	2, -1

- Defender(police)がTerminal 1を守ると、adversary(attacker)もそれを観察し、Terminal 2を攻撃する
- Defenderが最初の戦略をとり、それを知った上で Adversaryが戦略をとる→Stackelberg Equilibrium。
- 混合戦略を許すと、Defenderは、Terminal 1を1日の60%、Terminal 2を1日の40%守る、という表現も可能になる。この時、Adversaryにとっての不確かさが増大する。

図2-2 スタックルベルグゲーム

特に、利得行列において、Defenderにとって良い戦略がAdversaryにとって悪い戦略になっているようなスタックルベルグゲームのことをセキュリティゲームと呼んでいる。ロサンゼルス空港の警備、アフリカでの密猟者逮捕などが有名だが、さまざまな応用が広がっている。

自動交渉

自動交渉（図2-3）については、今日の発表3-2で詳しい紹介があるが、エージェント2者間の1対1交渉を扱う。交渉（Bargaining）は、もともと一つの財を二人で分け合う協力ゲームの枠組みがもとになっているが、非協力ゲームとして定義されることが多い。その交渉モデル（Bargaining Model）をRubinsteinが検討して^{[14][15]}、交互提案プロトコルを提案した。これは例えば、二人のプレーヤーAとBがいて、まずAからBへ提案し、Bがこの提案を受け入れるなら合意、受け入れないなら拒否して、逆にBからAに提案をするというふうに交互に提案をして進んでいくというモデルである。Rubinsteinは、このやり取りが無限に続くとしたら、理論的に均衡点が計算できるので、交渉する必要はないという結論を出した。しかし、現実的にはそうではなく、交渉に締め切りがあるので、それまでにどうやって合意するかという問題になる。それで、コンピューターシミュレーションが使われる。

そこで使われるモデルはいろいろあるが、その1つとして単調妥協プロトコルがある^[16]。エージェントは一つ前の提案より、相手にとって効用が低くなる提案をすることができないというモデルである。自動交渉の競技会では、そういったシミュレーターを用意して、エージェント同士が妥協して合意していくような仕組みを考えた。競技会では、妥協の仕方によって、より高い効用で合意をたくさん取るエージェントを見つけようとしている。

自動交渉 (Automated Negotiation)

計算機プログラム等を用いて**自動的に交渉の支援や代行を行うこと**。

自動交渉の利点

- ・ 複雑で大規模な交渉問題を合理的に合意形成できる。
- ・ 実際の交渉をシミュレートして解析できる。
- ・ **設計者の異なる人工知能間の自動調整**などに利用できる。

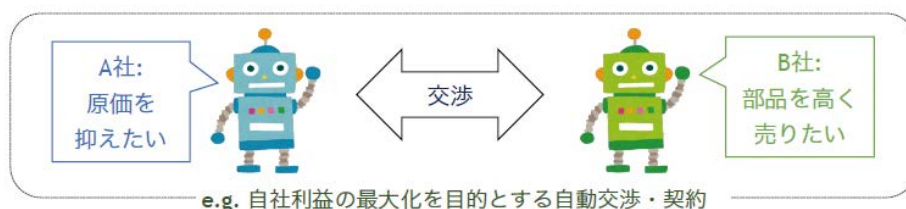


図2-3 自動交渉

集団知性・ハイパーデモクラシー

2000年前後からゲーム理論に基づいて発展したマルチエージェントが非協力的・競争的なエージェントになってしまったということの揺り戻しとして、もっと協力的な方向が大事だと言われるようになった^[17]。MITのThomas Maloneらによる2010年の論文^[18]では、人間同士の文殊の知恵、グループでの共同のパフォーマンスや集団知性といったものを、数値として測定することが試みられている。この論文では、グループに女性が含まれていると、その値が高くなることも示している。

ハイパーデモクラシープラットフォームの実現

ソフトウェアエージェントと人間が一緒に参加するSNS型の合意形成プラットフォーム

Augmenting SNS democracy Technologies

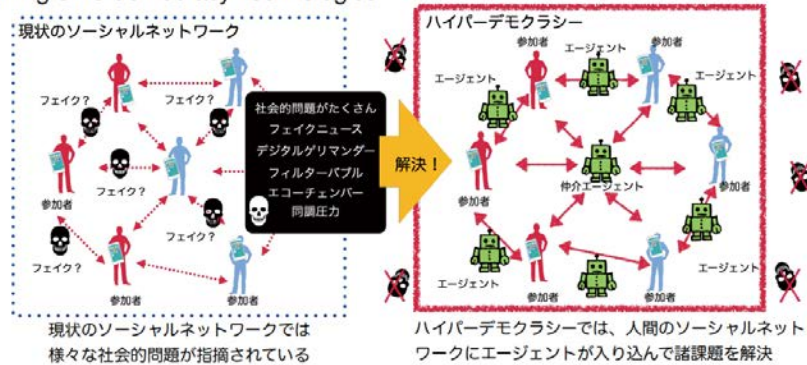


図2-4 ハイパーデモクラシープラットフォーム

そのような背景の下、「CREST 信頼される AI システムを支える基盤技術」の中で「ハイパーデモクラシー：ソーシャルマルチエージェントに基づく大規模合意形成プラットフォームの実現」を進めている（図2-4）。エージェントベースのプラットフォームで、議論・合意形成を支援したり、フェイクニュースやグリマンダリングを防いだりすることを目指している。

AIによって人間の知性を拡張するとか、身体性を拡張するという取り組みは多数あるが、さらに大事なこととして、社会システムや民主制といったものを、マルチエージェント技術で拡張して、より良いものにしていけることができていると思っている。

発表2-2：荒井 幸代

注目動向

私は、マルチエージェント系の意思決定規範をどう設計するかという問題に注目しており、技術的には強化学習（Reinforcement Learning：RL）や逆強化学習（Inverse Reinforcement Learning：IRL）に注目している。

関連動向を俯瞰すると、ルールベースから強化学習へ、単発行動後の結果、すなわち、一つ先の効用を扱うゲームから、逐次的な意思決定（Sequential Decision）を扱う確率的ゲーム、多目的逐次意思決定へとこの流れがある。実問題は行動の系列を扱い、行動の計画を立てる最適化問題の解決法として強化学習がアプローチの中心になっている。特に報酬に遅れを許すフレームワークなので、確率的ゲームと強化学習は非常に相性が良い。

また、ビッグデータが利用できるようになったことから、構成論から還元論へシフトし、逆強化学習、模倣学習（見まね）、多目的計画・多目的逐次計画などが注目される。ビッグデータによる設計論に持ち込みたいというのは、AAMASやAAAIなどのトップカンファレンスの傾向でも見られる。

対象問題とアプローチ

例えば、乗用車、バス、バイク、歩行者が混在した交差点での自動運転を考えてみると、いろいろ絡む中で、何を認識・注視するか、目的位置までどう行動するか、他車との競合をどう解消するかを考えねばならない。

また、築地から豊洲のような施設移転を考えてみると、周囲の交通網・建物なども含めて、何を認識・注視するか、どこに目標を設定するか、周囲への影響をどう予測するかを考えねばならない。

しかし、今のエージェント技術では、このようなことまではうまく解けていない。実は自動運転の例も施設移転の例も、同じような問題構造があって、Noisyな状況に非常に弱く、周囲やいろいろなものの関係性をうまく扱えていない。

対象問題とアプローチを図2-5にまとめた。このような問題をうまく定式化して取り組んでいく必要がある。解析・最適化については、深層学習、強化学習、進化計算など、有用な道具がたくさんある。関係性については、ゲーム理論、ネットワーク分析などあるが、まだ取り組みが少ないように思う。



図2-5 対象問題とアプローチ

エージェント規範の設計：マルチエージェント逆強化学習

最近の大きな変化としてビッグデータの活用がある。これを深層学習の教師データとして使うというのが多いが、逆強化学習によれば設計論にも使える。エージェントを設計する上で、「適正な評価規範」が要るわけだが、設計というと、昔はルールベースで、人がこうだろうという思い込みで作っていた。これがビッグデータから逆強化学習によって設計することができるようになった。

2つ目は「正確な状態認識」で、これについても状態ベクトルを定義するなど、以前は人の思い込みで作っていたが、ビッグデータから深層学習によって補完できるようになった。

3つ目は「最適行動の出力」で、これについては数理最適化という強力な方法があるが、モデルを必要とするという弱点がある。そこで、強化学習が活用されている。

4つ目は人が理解するための「説明の生成」だが、今回は省略する。

強化学習

強化学習（図2-6）は、逐次意思決定を扱う枠組みで、報酬というスカラー量を与えるだけで、試行錯誤によって最適行動を獲得できる。今の状態から次の行動を決めるワンショットなので、ゲーム理論であれば、今の状態から次の状態への利得が毎回与えられることになる。一方、強化学習の場合は、行動の系列に対して初めて与えられるものになる点が異なり、長期的な行動系列の計画問題に対しては強力だと考えられている。

強化学習では近年、報酬をどのようにして設計するかが課題となっており、主に2つの課題がある。 1点

目は、毎回行動してもなかなか報酬が得られないので学習しないという問題で、報酬のスパース性と呼ばれている。2点目は、報酬の定量化（スカラー化）が難しいという問題である。図2-7にその難しさを2つ示したが、定性的な感覚を定量化することが難しい上に、複数の目的を定量化することも難しい。

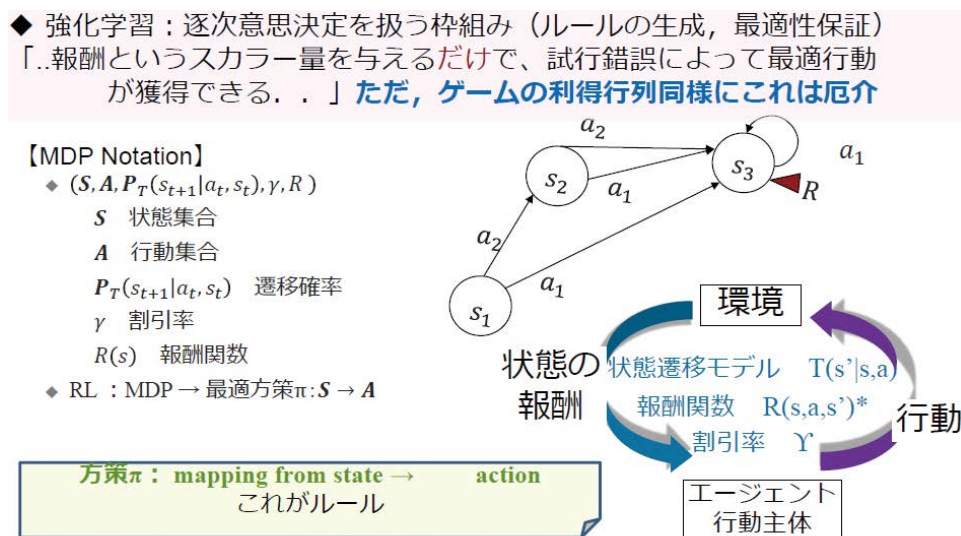


図2-6 強化学習

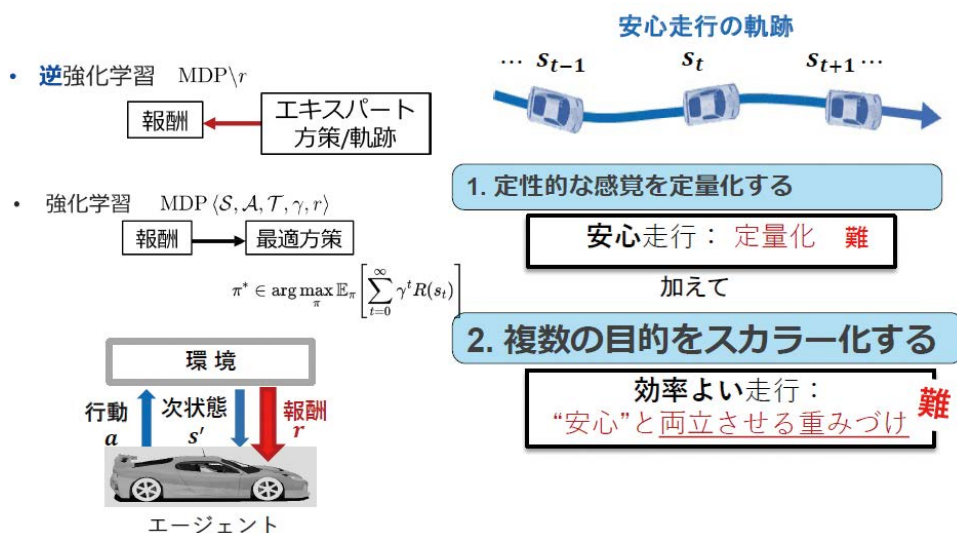


図2-7 強化学習における報酬の定量化の難しさ

模倣学習

模倣学習（Imitation Learning）は、教師あり学習や教師なし学習の限界への補償として注目される。AI研究者よりロボティクス研究者に設計論として注目されているように思う。1999年頃から取り組みがあり、

それが2016年頃から見直されてきた¹。今はもう相当熟してきた時期だが、日本での取り組みはやや少ない²。

ここで、模倣学習と Behavior Cloning との違いを説明しておきたい。Behavior Cloning は、エキスパートや過去にうまくいった実績から、その動作・行動をコピーするというものだ。それに対して、一般的な模倣学習は、1回コピーした上で、そこから改めて自分で報酬関数を推定して、自分なりの方策を作り直す。これによって結果にどんな違いが生じるかというと、Behavior Cloning ではコピーしたものが問題なく動いているうちはよいが、実行環境が少し変化して失敗してしまったとき、挽回しようがない。しかし、模倣学習では、もともとのエキスパートや過去の実績では経験されていないような状況や異なる環境に対しても頑健になる。

逆強化学習

その報酬関数の推定に使われるのが逆強化学習（IRL）である。図2-8に示すように、強化学習（RL）は、報酬が与えられて、最適行動を決定する。それに対して、逆強化学習（IRL）は、最適行動が与えられて、報酬関数を決定する。ここでいう最適行動は、先ほどの模倣学習の説明では、エキスパートの意思決定の系列、行動の軌跡が該当する。

逆強化学習の用途としては、エキスパート方策の模倣という面では、ロボットアームの制御や自動運転などにも適用されている。実機を動かす際に、事前のシミュレーションで想定していた環境から実機環境が多少変わっても何とかするという頑健性を期待されている。また、人間や動物の行動解析にも使われていて、例えば、歩行者の行動解析、SNSユーザーのトピック変化解析など、人間が何をモチベーションとして動くのかを推定するといったことが試みられている。これはマルチエージェントシミュレーションにおけるエージェントの設

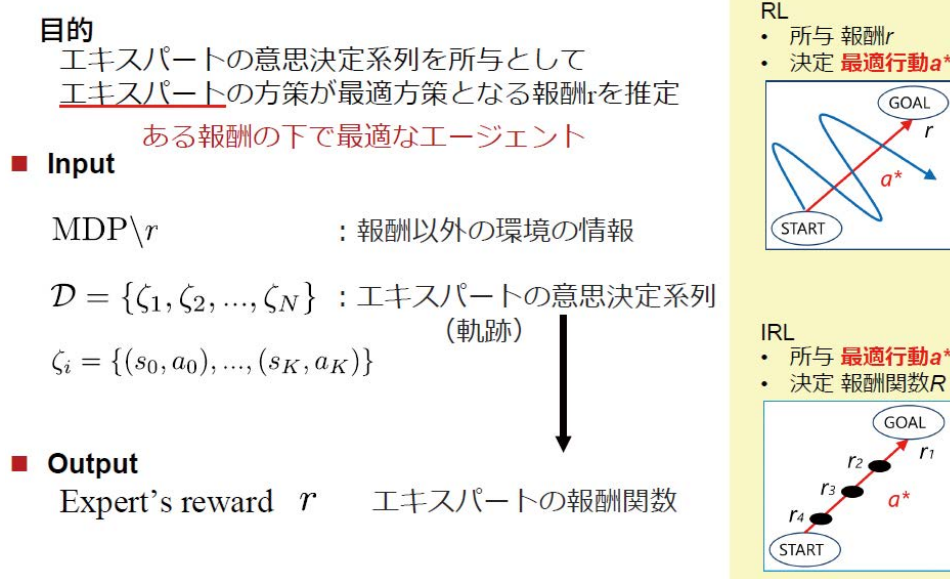


図2-8 逆強化学習

1 2018 ICML Tutorial :
<https://sites.google.com/view/icml2018-imitation-learning/>
 2016 ICRA Tutorial :
<https://www.icra2016.org/conference/tutorials/>

2 日本では2000年前後から認知発達ロボティクスの構成論アプローチがある（浅田先生、國吉先生、尾形先生など）

計にも使える。不審者や異常な行動の検出といった利用例もある。

ここまで、基本的には逆強化学習単体の話だったが、もちろんマルチエージェント系にも拡張できる。2体系とN体系など、マルチエージェント逆強化学習の分類と主要論文を付録2に載せておく^[19]

今後の課題・展望

人間の行動をエージェントに投影する必要性が高まっている中で、ビッグデータを活用できることから逆強化学習に注目した。これから特に重要なのは制御理論との連携だと思う。実際、制御系の研究者とモデルベースやモデルフリーなど、精力的に話をしている。このあたりの最新動向について、国際会議としては、AAMAS (International Conference on Autonomous Agents and Multiagent Systems) よりも、むしろ IROS (International Conference on Intelligent Robots and Systems) の方が活発に取り組まれているので、見ておく必要があると思う。

それからもう一つ、投影できても、いかに全体最適にするか、というのも重要課題である。例えば、エージェントAはこちらへ行きたい、エージェントBはあちらへ行きたいというようなときに、現状、パレート最適が複数あるというところで終わってしまっている。それに対して、ここに行きたいというそれぞれのエージェントに、どのような追加報酬 (Reward Advancement) を与えなければならないかという問題が割とホットになっている。均衡解から最適解への制御ということで、多目的強化学習 (Multi-objective RL) や多目的逆強化学習 (Multi-objective IRL) が重要になってくると思っている。均衡から全体最適へというのは昔から言われていたが、その具体的な方法はまだ確立されていない。このニーズが強いのはロボット系で、Noisyな環境で収束させたいということも含めて、リアルな問題になっている。これは社会系のAIでも、本当に解決しないといけない問題ではないかと思っている。

討議

青木 伊藤先生が最後のあたりで話された、エージェント技術で民主主義を研究するという点について、もう少し知りたい。

伊藤 インターネット上のフェイクニュースやデジタルグリマンダリングなど、民主主義をうまく機能させないような問題に対して、エージェントによってそれを防ぐようなプラットフォームを構築したいというのがモチベーションである。その中に、マルチエージェントシステムで培われたゲーム理論や、ソーシャルチョイスのためのボートティングの仕組みが使えると思っている。今CRESTで取り組んでいる。

山田 荒井先生が最後に触れた多目的逆強化学習は、マルチエージェントがそれぞれ報酬関数を持っていて、そのパレート最適を逆強化学習するということか。

荒井 その通り。今のところは各エージェントがある目的関数1個で動いていて、それが複数あれば複数の目的になる。すると、そのN人のエージェントのパレートフロントが出てくる。しかし、今はどれか1個の均衡に決めなくてはならないというのが現実である。そのとき、各エージェントはどのパレートフロントでもいいというわけではないので、1個のパレートフロントの中の解を選ぶための追加報酬を設計しようという方法論に向かっている。また、各エージェントが複数の目標を持っている場合に拡張もできる。

伊藤 その話は、我々が取り組んでいる自動交渉ととても問題意識が似ていると思う。

荒井 私もそう思った。自動交渉は、ある行動ごとにやるのか、それとも、あるシーケンスでやるのか。プランとプランの比較もやるのか。

伊藤 逐次意思決定に近い。

森永 私が自動交渉について発表するので (発表3-2)、そこで関連する話ができると思う。

ただ、プランの調整の意味を確認したい。例えば、移動体の軌道計画で、自分はここを通るから、あなたはこの時間帯はあちらにいてほしいというような交渉をすることなのか、それとも、交渉しながら効用関数を変えていくような話なのか。前者はまさに我々が取り組んでいるし、後者は一部対応している。

荒井 自分は前者を考えていた。全体最適になるような解をパレートフロントで見つけられたとしたら、前者の問題は役割分担、効用関数の重みの変更という形で考えている。後者の方は今のところ対応していない。

福島 環境・状況がどんどん変わっていくようなケースで、強化学習はそれなりに追従していけるのか。それとも、試行錯誤を通してできるのは、あくまで今の時点のモデルと考えるべきもので、追従は難しいか。

荒井 それに関連して、リアルタイム強化学習がホットピックになっている。その変化の時間がどれくらいなのかにもよるが、過去事例から方策を複数通りため込んでおいて切り替えるという方法や、リアルタイムに追従するという方法が出てきている。常にサンプリングしながら環境の変動を検出しようという考え方がある。それをマルチエージェントでやる手もある。

福島 多目的のマルチエージェントというとき、各エージェントがそれぞれ違う効用関数を持っていて、自分の効用関数を変えていくというようなことも考えられているか。

荒井 まさに私のところで検討している問題に近い。現状、ワンショットに対して効用関数を決めているので、シーケンスになると、シーケンス中のどこでも効用関数の重みは同じということになる。しかし、最初は効率性重視だったけれど、安全性重視に切り替えるというような、効用関数の重みを変えていくという考えを入れないと、長いプランの場合、逐次意思決定においてパレートフロント解を多数見落とすことになってしまうことが、実験的に分かってきた。もっと良いフロントパレート解集合があるはずで、候補が多いほど合意が得やすいことになる。そういう研究はまだほとんど行われていない。

福島 なるほど、そういう切り替えは重要な取り組みだと思う。さらに難しいことを言うと、人間同士の交渉では、欧米的にはディベートで勝ち負けを決めるような感があるが、日本人的な交渉だと、相手に共感して自分の効用関数を変えたり、止揚（アウフヘーベン）によって上位解を見つけたりといったことをすると思う。そういうレベルのことまでエージェントができればすごいことだが、簡単ではないと理解した。

山田 効用関数がダイナミックに変わるというのは、概念的・研究的にはありだと思うし、現実でも起きることだが、その問題設定だと定式化や学習がはるかに難しくなるのではないかな。解ける方法がほぼないというようなことにならないか。

荒井 効用関数を目的の重み付き線型結合のようなシンプルな形で表現するのであれば、それほど難しくはない。実際、ベンチマークでそれぞれの効用の重みの変化を見つけることは一応できていて、理論というよりも、計算量が問題になっている。

3 | マルチエージェント領域（2）

発表3-1：野田 五十樹

社会制度設計のためのマルチエージェント社会シミュレーション

最近、私が重要だと思っていることとして、多数のステークホルダー間の調整を挙げたい。これが、社会的にも、人工知能の研究としても大切だと思っている。

これまでの人工知能の研究では、最適化や効率化が中心になっていて、論文も性能が数パーセント向上したというようなものが多い。しかし、本当に大事なのは社会制度設計に応用していくことで、そのツールとしてマルチエージェント社会シミュレーションMASS（Multi-Agent Social Simulation）が重要で、そのためのシミュレーション・ビッグデータマイニングやHPC技術も充実させていく必要がある。

MASSは、対象が社会システムなので、普通のシミュレーションに比べると結構ほわっとしているというのが、難しいとともに面白いところだ。モデルが決まってない、ステークホルダーや目的関数が多い、境界条件が曖昧といったことから、物理実験のように問題を簡単に切り出せないことが多い。ネットワーク外部性もあって、大規模なシミュレーションが必要になるという点も難しいが、研究しがいがある。ビッグデータが取れても、不完全なところがあり、それをどう補うのかもポイントになる。

多数のステークホルダー間の調整に向けた取り組みとしては、公平性などを考慮したゲーム理論的アプローチや、長期の時間発展の予見を可能とするシミュレーションや、適応や行動変容を陽に表現できるエージェントモデリングなどがある。私自身はどちらかというと、力技で網羅的にシミュレーションを回すことで、多面的な分析・最適化を試みるアプローチを重視している。

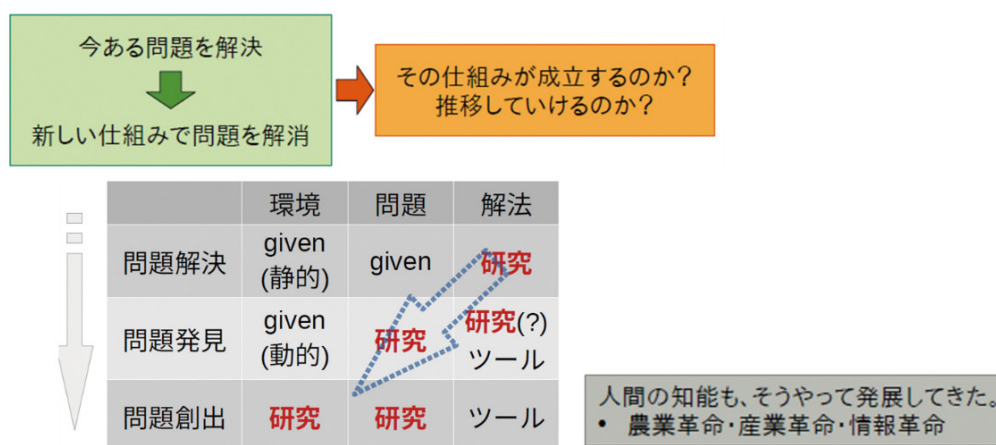


図3-1 問題解決から問題創出へ

問題解決から問題創出へ

少しメタな話になるが、科学研究、特に工学は問題解決を狙って、その解法を見つけるというのが、普通は取り組みの中心になる。しかし、人工知能や社会シミュレーションでは、問題を新たに発見するとか、問題を創り出すといった方向へ研究がシフトしてきている（図3-1）。

こういう方向で考えると、新しい仕組みになれば、今問題になっていることがそもそも存在しなくなって解決することになる。そのとき、そのような新しい仕組みの社会に移っていきけるのかとか、それが維持可能なのかを考えることが重要になる。

問題創出の例(1) オンデマンド型公共交通

少し手前みそになるが、我々のところでSAVS（Smart Access Vehicle Service）と名付けたオンデマンド型公共交通サービス、要はAIによってタクシーの乗り合いマッチングを行うシステムを事業化している。これは交通というものを考え直そうというところから話が始まり、まず乗り合い型が成立するのかというところから研究を始めた。それで採算が取れると見込めたところで、次に、よりいろいろなことと結びつける道具として、こういった交通が使えないかと考えて、今でいうMaaS（Mobility as a Service）的なコンセプトの走りに2000年代から取り組むことができた。最近ではMaaSの考え方が普通になり、そうすると、最適化しようという研究ばかり増えているが、むしろ重要になってくるのは、そこでたくさんのステークホルダー間で調整するところだと思う（図3-2）。

●何が課題なのか？

- ▶ 日本型MaaSで必須となるステークホルダー間の**連携**のための**調整の枠組み**の欠落
 - 多くの事例で、採算性の問題からMaaSの持続性が懸念。
 - 採算性確保のための処方箋はサービス・ステークホルダー間の連携と効率化。
 - 連携を成立させる最重要課題は、連携参加を促せる**公平な調整機構**

「日本型」
 多様なプレーヤーで構成されるMaaS
 ユーロッパ型（行政運用主体）、米国型（プラットフォーム独占）とは異なる条件

●日本型MaaSを成立させるため3つの技術フェーズ。

- ▶ **基礎理論**：ステークホルダー間の対立構造解消のための**調整技術・手法**の確立
- ▶ **実用化**：実運用レベルに近い条件での**マルチエージェント社会シミュレーション**により、基礎理論の動作・機能を確認。
- ▶ **社会実装**：シミュレーションによる実用化サービスを現場実装し、運用結果を**分析・評価・モデル化**。基礎理論や実用化フェーズへフィードバック。



図3-2 日本型MaaS実現に向けて

調整に関する研究も行っていて、例えば、配車をするときに、タクシー運転手の収入の平等性も考慮しながら配車をするというアルゴリズムを開発し、パフォーマンスをあまり落とさずに収入格差を改善する方法があり得ることを、シミュレーションで示した。また、いくつかのタイプの運行スタイル、異なる価格設定を持っているタクシー会社が競争しながらサービスを提供している状況のシミュレーションも行った。利用者も学習をするし、会社の側もGA的にいろいろなビジネスモデルを試すようにしたら、興味深い現象が生じた。最初は各社、価格競争をするが、その後、ある時点から急に「ぼったくり」が出てきた。このことから、「ぼったくり」を防ぐ制度を最初から設計しておかないと、自由競争がうまく成立しないということが示唆された。シミュレーションの結果から他にもいろいろ示唆が得られている。

問題創出の例(2)COVID-19 AI・シミュレーションプロジェクト

もう一つのホットピックとして、コロナ対策のために取り組んでいる内閣官房 COVID-19 AI・シミュレーションプロジェクトがある。現在は感染シミュレーションが中心だが、今後は医療リソースとのバランスや社会活動とのバランスなどにも広げていくべきと思っている。図3-3は倉橋先生による感染シミュレーションの事例だが、倉橋先生は社会活動とバランスを取ることを意識したシミュレーションにも取り組んでいる。

●エージェントの街を仮想し、感染拡大を模擬

- ▶ 東京近郊として現実的な家族構成
- ▶ 外出頻度と会食人数を変化させ、患者数の変化分析

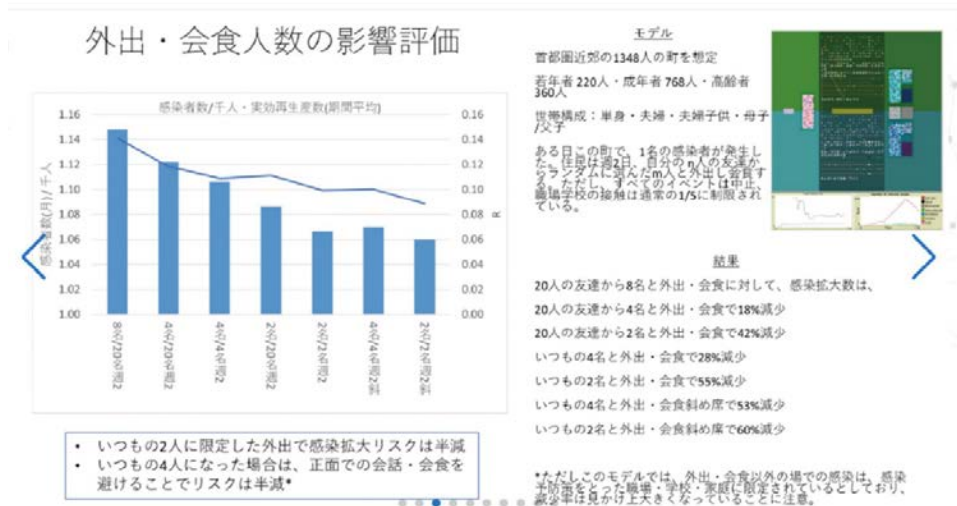


図3-3 仮想的な街での感染シミュレーション³

マルチエージェント社会シミュレーションの応用分野・課題

マルチエージェント社会シミュレーションの応用分野として、公共交通と疫病対策について紹介したが、それ以外にも、都市・地域設計や防災・減災、電力市場・金融市場・物流市場といった市場設計など、いろいろ取り組まれている。

このような応用では、膨大な計算が必要になる。スパコンも活用しているが、富岳や京コンピューターはマルチエージェントシミュレーションの大量実行にはあまり向いてなくて、産業技術総合研究所のABCIが今一番使いやすい。エージェントシミュレーションに向けたスパコンやHPCのアーキテクチャーや運用も重要な課題だと思う。

また、今後、サービス連携がますます重要になってくる。スマートシティというと、交通がクローズアップされることが多いが、交通を利用しながら、それを要としてさまざまなサービスを連携させる場と考えた方がよい。そのとき、マルチステークホルダーの間でどのように利益配分やコスト負担をするかがポイントになる。それを成立させる条件をきちんと分析することが望まれて、それを可能にするのがマルチエージェント社会シミュレーションだと思っている。

3 COVID-19 AI・シミュレーションプロジェクト：
<https://www.covid19-ai.jp/ja-jp/>
 同プロジェクトにおける倉橋先生のご取り組み：
<https://www.covid19-ai.jp/ja-jp/researcher/setsuya-kurahashi/>

発表3-2：森永 聡

自動交渉技術の必要性

企業の立場から、自動交渉技術の研究開発と実用化・社会実装を中心に話し、最後に、エージェント関連で注目している他の技術についても触れたい。

スマート化されたシステムが広く普及して、十分に社会価値を発揮するためには、それらの間の挙動や利害を調整する機能がキーになる。例えば、自動車、ドローン、潜水艦といった自律移動体は一つ一つがスマートシステムだが、実空間上で出会ったときに、軌道計画や経路計画を調整できないと、お見合いを起こしたり、急いでいて周りはどう動くか分からない中、突っ込んでしまったりする（図3-4左）。

- ◆ スマート化されたシステムが広く普及し十分に社会価値を発揮するためには、それらの間の挙動や利害を調整する機能がキーになる → 個別最適化・内部最適化の限界を突破
- ◆ 既存のデータ共有や協調制御のアーキテクチャは、参加者の内部情報開示や自己決定権剥奪を前提としている → 経済主体間の調整では、相談や交渉ベースのしくみが必要/非常に有益

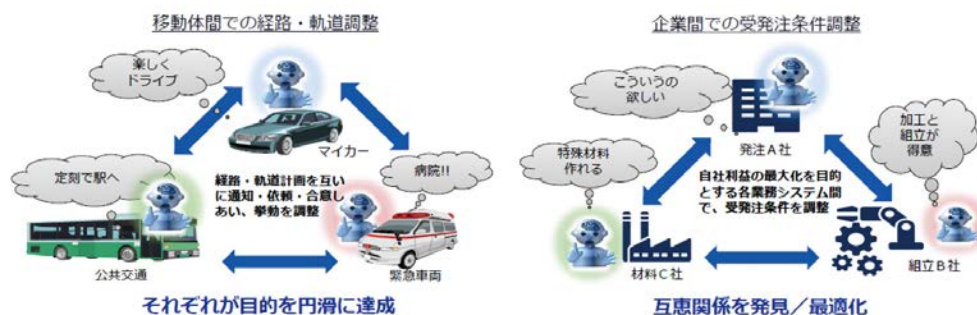


図3-4 スマートシステム間の挙動／利害調整

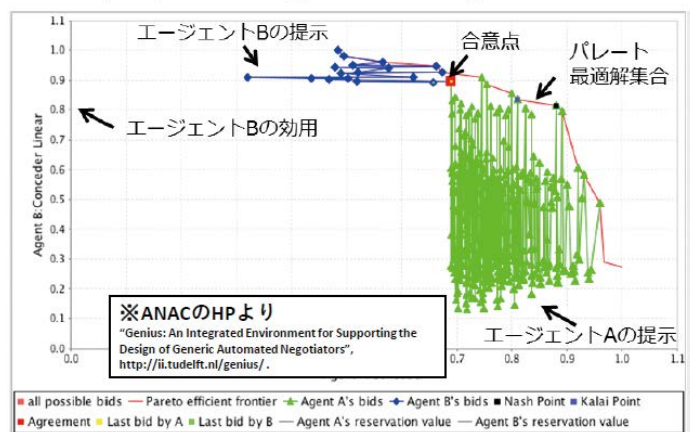
また、現在は会社の中を効率化するスマート化のソリューションが提供されているが、それは要するに、会社の中だけ効率化しているものだ。しかし、会社と会社の間で、納期などの受発注条件を調整できたら、無理に残業する必要がなくなって、もっと良い結果が得られるかもしれない（図3-4右）。

つまり、現状の効率化・スマート化は、実は個別最適化・内部最適化になってしまっていて、それが意味で限界になっている。外部と調整できれば、その限界を突破し得る。

ただ、既存のデータ共有や協調制御のアーキテクチャーは、参加者の内部情報開示や自己決定権剥脱を前提としている。それができる場合はそれでよくて、例えば、ディズニーランドの中のカートの群れの制御とか、社会主義国における計画経済とか、1つの会社のコングロマリットの中のリソース融通とかが該当する。NECでは、そういった場合の仕組みについても研究開発しているが、それができないケースもあって、そのようなときに必要になる仕組みについても頑張っており取り組んでいる。その最たるケースが経済主体間の調整で、このケースでは相談・交渉ベースの仕組み、自動交渉技術が、必要かつ非常に有用なものになる。

自己利益（効用関数）の最大化を目的とした交渉エージェントの実現

- ・相手に提示すべき合意条件案の自動生成
- ・提示された合意条件案に対する受諾/拒否の自動判断



国際的な競技会 Automated Negotiation Agent Competition (ANAC)が毎年開催されている。

図3-5 相手との相談・交渉を支援・自動化するAI技術

自動交渉エージェントの技術

交渉エージェントでは、自己利益を効用関数の形で表現する。図3-5は、横軸がエージェントAの効用、縦軸がエージェントBの効用で、交渉が進むにつれて効用関数の値が推移していく様子を表している。エージェントは互いに相手の効用関数を知らない。最初のうちは妥協したくないので自分の効用が高くなる提案ばかりし合うが、時間が経つにつれて、双方の妥協可能ラインが下がってきて合意点に至る。

先ほどのセッションで、合意に至る過程で効用関数を変化させることはないのかという議論があったが、この図3-5のケースでは効用関数は変化せず、定常なものとしている。しかし、非定常な効用関数についての研究も行っている。例えば、救急車の進路交渉をするエージェントがいて、もともとは安全性重視としていたが、乗せている患者の容体が急変したら、スピード重視に切り替えるというような、合意点を見つける前に、短いスパンで効用関数を変えるというのも、実用上、絶対必要だと考えている。

自動交渉の応用事例

自動交渉の応用として取り組んでいる事例をいくつか紹介する。

まず、移動体間の経路・軌道調整の実用化に取り組んでいる。自動車は込み入った社会受容性の問題があるので、まずドローンのケースに取り組んでいる。デモシステムは地上移動のケースも試作した。例えば、警備サービスA社、掃除サービスB社、運送サービスC社それぞれが、AGV（Automatic Guided Vehicle：無人搬送車）を動かしている状況を考える。それぞれは自分の目的で動いているので、それらが混在すると、経路の調整や計画の調整をしないと、衝突したり、動きが取れなくなったりする。事前に経路・計画を全て決めておくとか、中央管制から全てに細かく指示を出すとかは、非現実的、ほぼ不可能だけれど、自動交渉ベースでやれば問題を回避できる。実際に、NEDOのプロジェクトの中で、稚内で別々の業者が飛ばすドローンの間の経路調整が自動交渉ベースでうまくいくという実証実験を行い、プレスリリースも行った⁴。

4 NEDO「ロボット・ドローンが活躍する省エネルギー社会の実現プロジェクト／地域特性・拡張性を考慮した運航管理システムの実証事業」で実施し、プレスリリースを行った。
https://jpn.nec.com/press/202111/20211104_03.html

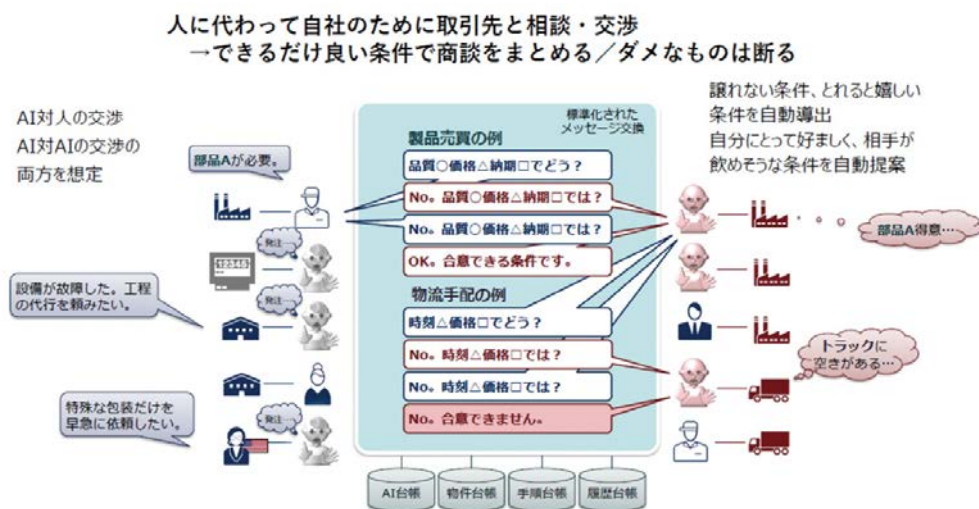


図3-6 企業間での取引条件の調整

別の応用事例として、会社と会社の間の商取引条件の調整に自動交渉を適用している。物を売りたい会社、物を買いたい会社、物を運んでほしい会社、物を運びたい会社などがあるとき、その間の商取引条件は、通常、セールスマンや調達担当者が必ずいて、言葉で交渉している。ここに、自分の会社のために取引先と交渉・相談をしてくれるエージェントを開発・導入する。すると、人に代わりにエージェントが、取引先と折り合える商取引条件をメッセージ交換で探ってくれる（図3-6）。自社も取引先も交渉エージェントを使っていると、人間よりもはるかに高速にメッセージをやり取りして交渉を済ませてくれる。

また、製造分野でのエージェント間の自動交渉も試作した。日本では、サプライチェーンが非常に深く長く、多数の会社間調整がいろいろなところで行われている。1つの会社を見ると、同時に複数の相手と商談を進めているとき、それらの商談が整合しているかが問題になる。例えば、物を売る約束をしたとき、それに必要な部品が調達されていないといけなないので、複数の商談を整合的に取りまとめることをAIで行えるようにもした。

移動体間の調整で実証実験を進めているほか、製造・商取引の応用ではシステムを見本市での展示⁵もしている。

社会実装のための活動

このようなマルチエージェント自動交渉の社会実装では、交渉相手のエージェントや人間との間で、交渉メッセージが正しく解釈されないと話にならない。そこで、交渉メッセージ交換に関する国際標準化を、国連の下で標準化団体（UN/CEFACT）にeNegotiationプロジェクトを提案して採択された。2022年上期に標準公開の予定で、パブリックコメントを受け付けて、ほぼ大丈夫というような段階にある。

このような取り組みは、広めていけばいくほどサプライチェーン全体として効率が上がるので、普及のためにいろいろな業界団体に働きかけている。具体的には、国際業界団体での公式テストベッド化を目指して、Industrial IoT ConsortiumにNegotiation Automation Platform Testbedを提案して採択され、日独韓からなる7機関体制で活動している。

また、自動交渉の国際競技会としてANAC（Automated Negotiation Agent Competition）が開催されている。いろいろなルールでもう10年以上実施されているが、2019年から商取引のサプライチェーンでの

⁵ 国際物流総合展2021、スマート工場EXPO2021など。

交渉技術を競うSCMリーグを立ち上げた。出場者は、サプライチェーンの中の1つの会社として、部品を買う交渉を行い、製造計画を立て、製造した製品を売る交渉を行うという3つを、上手に整合的に実行する交渉AIを提出する。仮想経済空間内でそれらを動かし、一定期間で最大利益を獲得した交渉AIが勝者となる。初年度（2019年）8チーム、翌年度（2020年）22チーム、昨年（2021年）51チームと出場者が増加し、大変盛り上がっている。

もう一つ、自律調整SCMコンソーシアム⁶を立ち上げた。想定ユーザー、想定ベンダーを組織化して、大学の先生にも参画していただき、自動交渉の社会実装を推進している。

サプライチェーンにおいて日々発生している「企業・組織・個人間での利害や挙動の調整業務」を劇的に効率化することを目指し、先進技術を活用した実用的な調整業務フローの整理と検証、その発展と普及を行う。

現在、一般会員8機関、ユーザー会員15機関、学術専門会員10名/機関

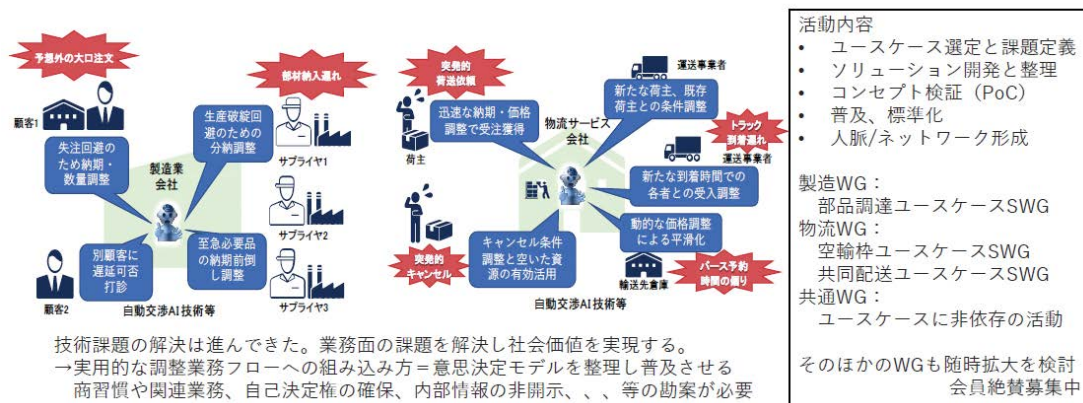


図3-7 自律調整SCMコンソーシアム

自動交渉技術は一定のレベルまで進んできた。これを実用化へ進めるには、今の商習慣や周りの関連業務と整合させる必要があるとか、日本では計画経済はできないので自己決定権は確保するとか、商談に必要な情報以外のいづれぐらい在庫が余っているかといった内部情報は絶対に非開示とするとか、そういった要件を全て満たすような自動交渉技術が期待されている。

このような先進的な取り組みをしているのは、まだ我々だけだと思っていたら、2019年創業のシリコンバレーのベンチャー Pactum 社が、いきなり2020年にウォルマートで試験運用を始めた^[1]。2021年の公開情報では、ウォルマートの調達部門はこのPactum社の自動交渉ソリューションを使ってサプライヤーと交渉していて、Maerskという世界最大の海運コングロマリットの調達部門もそれを使って陸上輸送に関する調達を行っていて、電気電子機器のWesco Internationalも何か使っているらしい。ドイツでもマイサイプライという会社が自動交渉ソリューションを販売開始している。

自動交渉以外で注目しているエージェント関連研究動向

自分が注力しているのは自動交渉だが、それ以外に会社として注目しているエージェント関連研究動向としてはゲームAIがある。スクウェアエニックスの三宅陽一郎氏に注目しているというのにほぼ等しいのだが

6 自律調整SCMコンソーシアム：
<https://automated-negotiation.org/>
プレスリリース：
https://jpn.nec.com/press/202109/20210917_01.html

[2][3][4]、彼が提唱している MCS-AI 動的連携モデルは、キャラクター AI、スペーシャル AI、メタ AI がいて、全体としてパレート最適解のどこに落としどころをつけるかを考えている。

それから、もっと広い意味でメタバースにも注目している。Facebook のメタバースだけでなく、NVIDIA の Omniverse や、Tencent や Microsoft による関連買収の動きとかがあり、別々の人が作ったエージェントが仮想世界で経済活動する時代のエコシステムは、エージェント関連から見ても気になる動向だと思っている。

討議

福島 交渉のテクニックとして、ブラフを使ったり、相手を騙したりといった手段もある。自動交渉の競技会では、そのようなテクニックも使われているのか、それとも、禁止やガードがかけられているのか。

森永 競技会の中では、心のきれいなエージェントしか参加できないトラックと、何やってもいいトラックがある。SCM リーグでは、自分が提出したエージェントは、仮想経済空間の中で 1 つの会社しか担当できないことにしているトラックと、複数の会社を担当できるトラックを設けている。後者は談合トラックと呼んでいて、談合のようなことをやってもよいことにしている。前者の方はそれを禁じている。ブラフはちゃんとした交渉テクニックなので、悪いことではない。普通の経済空間の感覚でいえば談合は悪いことだ。また、意図せずに変なことになるという現象もあることが分かって、東大の和泉先生と一緒に研究している。例えば、似たようなロジックを持つ交渉エージェントが経済空間内にいくつものいると、思わぬ方向の暴落や高騰が起きることが確認されている。それから、エージェントが強化学習で自分の交渉テクニックを磨いていくと、自然に談合を始めるという現象も知られている。

福島 マルチエージェント社会シミュレーションを進めていく上で、人間社会の現象でエージェントのモデルに寄せにくいところとか、どうしても外れてしまうところとか、現状、そういった難しさはどのあたりに残っているか。

野田 まず、取得できたビッグデータから、ある程度エイヤとモデルを作らざるを得ない。交通でいえば、人々が何時頃動きたがるとか、どういうところに行きたがるかといった傾向は、ある程度分かるが、それでは捉え切れていないことは当然たくさんある。現状は、乱数を入れて、わざと外れ値を出してシミュレーションに乗せることを繰り返して、起こり得る可能性を広げるようなことをしている。しかし、現実との乖離は出てくるので、そこをきちんと分析することで、モデルをより精細化できる。ただ、それで全部をカバーし、説明し尽くせるわけではない。

福島 マルチエージェント社会シミュレーションは、正確に予測したり、予測的中率を上げたりということよりも、こんなことまで起こり得るという、さまざまなケースを事前に把握・想定できることがポイントだと思う。人間が考えたときよりも幅広くさまざまなケースを想定できて、その上で、戦略や対策をプランニングできるというのが、現状、大きなメリットなのだと思うが、この理解は正しいか。

野田 その理解でよいと思う。最近のコロナ感染シミュレーションで、研究者は推移予測として幅のあるグラフを出しているのだが、なぜかマスコミからは 1 本線になって出てくることがよくあり、結構問題だと思っている。起こり得る可能性の幅をきちんと捉えた結果を出すことは、社会シミュレーションでは必須だと思う。

山田 マルチエージェントシミュレーションで人間の行動をシミュレートするとき、行動経済学でいう認知バイアスの・非合理的な振る舞いは、どれくらい取り込まれているか。また、人間の集団のヘテロジェニアスな部分はどれくらい取り込まれているか。

野田 認知バイアスは当然、入れることは考えている。例えば、限定合理性のような、認知可能部分を限定することで表現できたり、時間的に揺らぎを与えたりすることはできる。ただ、それを入れたとき、他

の部分の揺らぎもあって、結果がよく分からなくなるところがあるので、簡単ではない。ヘテロについても、例えば、時間に追われている人とそうでない人といったバリエーションを入れる試みはしているが、それをサポートするデータがなかなか手に入らないといった問題と合わせて考えないといけない。また、アンケートを取るという試みもある。実際に使った人とか使うつもりがある人に、何が重要かを選んでもらって、交通選択モデルを作るようなことは試みている。そこから発展させられないかと考えている。

山田 ヘテロも認知バイアスも、とりあえず無視しても十分な性能は出るということなのか、やはりそれを入れないと駄目なのか、どちらかと思っているか。

野田 最終的には入れなくてはいけないと思っている。ただ、今はサービスが成り立つことを、行政の人とかに納得してもらうといった段階で苦労していて、もっと手前の課題がいろいろある。また、ヘテロを入れることによって何が明らかになるのか、研究テーマの面から詰めていく必要もある。

森永 自動交渉では、ヘテロ性と揺らぎ、あるいは、非合理的な交渉相手といった点は、絶対考えないといけない問題だと思っている。交渉相手のタイプを見極めて、こちらの出方を変えないと失敗してしまう。第1回の競技会で優勝したエージェントは、下手な取引相手を見つけて、そこから徹底的に搾取するという方法を使ったので、ヘテロ性への対応に成功して優勝したと言える。

山田 人間はその場で相手のモデルを作りながら交渉していると思うが、自動交渉はそういった方向に向かうのか。

森永 それに加えて転移学習もポイントになる。ゼロから対戦相手との会話で学ぶというよりは、よくあるタイプのどれに相手が属するかを考える。これに対して、競技会で優勝したあるエージェントは、相手は自分がどういうタイプか推定しようとするのを見越して、自分がどういうタイプなのか、だまし続けることで優勝した。

4 | HAI（Human-Agent Interaction）領域

発表4-1：山田 誠二

HAIとは

HAIは日本発の研究分野で、2005年頃に、私と、次に話をする小野哲雄先生、それから今井倫太先生、竹内勇剛先生、岡夏樹先生あたりが中心になって立ち上げた。HAIをなぜ始めたかという、個人的には、マルチエージェントシステムに少し不満を感じていた。当時のマルチエージェントシステムは、基本的にホモジニアスで、人間がシステムに入っていなかった。また、最終的にAIを使うのは人間なのに、人間に使いやすいAIという観点がないといったところにも疑問を感じた。国内でHAIシンポジウムを始めた後、HAIの国際会議も始めて今年で10回目となり、世界的にも認知されてきたと思っている。

エージェントの定義として、通常は「自律的に行動する主体」のように捉えられているが、HAIでは「私たちが人間のように感じる（擬人化する）人工物」を全てエージェントと呼ぶ。ある意味で、非常に主観的に捉えている。基本的にMedia Equationの考え方で、自然物を対象に入れてもいいのだが、HAIは工学の分野なので、対象を人工物としている。

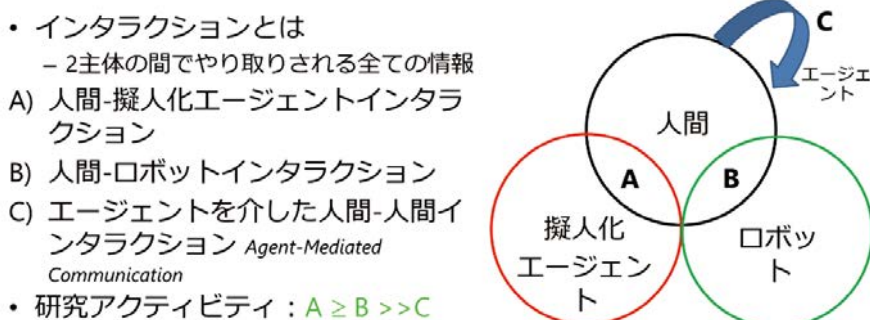


図4-1 HAIにおけるインタラクション

エージェントの種類については、擬人化エージェント（仮想エージェント、チャットボット）、物理的身体を持ったロボット、人間そのものという3種類がある。また、HAIにおけるインタラクションは、人間と擬人化エージェントの間のインタラクションA、人間とロボットの間のインタラクションB、人間と人間のインタラクションCという3種類がある（図4-1）。ただし、Cは、人間と人間のインタラクションというと社会科学的な話になってしまうので、ここでは、エージェントを介した人間と人間のインタラクション（Agent-Mediated Communicationと呼ばれる）に限定している。インタラクションの定義も難しいが、HAIでは「ある2つの主体の間でやり取りされる全ての情報」としている。例えば、外見も、相手にとって視覚から入ってくる一種の情報なので、インタラクションの一つとして扱う。それで、これらA、B、Cという3つのインタラクションをどのように設計するかがHAIの研究対象になる。

実際の研究アクティビティとしては、ほとんどの研究が人間と擬人化エージェントの間のインタラクションA、あるいは、人間とロボットの間のインタラクションBに取り組んでいる。人間と人間のエージェントを介したインタラクションCは、今のところ非常に少ない。

HAI デザインの目的と対象

HAIは、人間とうまく協調できるようなエージェントを作ることとしたインタラクションデザインの技術開発である。つまり、先ほど述べたように、やり取りされる全ての情報をいかに設計するかということだ。それによって、人間とエージェントがうまく付き合えるとか、エージェントが人間の良きパートナーになれるといったことが、HAI 設計の目標である。

では、具体的には何をデザインするのだが、その1点目はエージェント自身をどうデザインするか。これは、エージェントの動きをどのように決めるかとか、どのような学習機能を持たせるかといったアルゴリズム開発が含まれる。

2点目は人間とエージェントの間でやり取りされる情報をどうデザインするか。これは典型的には、エージェントの外見のデザインや、人間に対してどんな情報をどんな表現、モダリティーで表出すれば良いかというデザインのことである。

3点目は人間とエージェントの関係をどうデザインするか。例えば、協調作業のときに、人間とエージェントがどのような社会的関係を構築すべきか、どのようなタスクを持ってくればいいのか、どのような種類のゲームと一緒にやればいいのかといったことを考える。

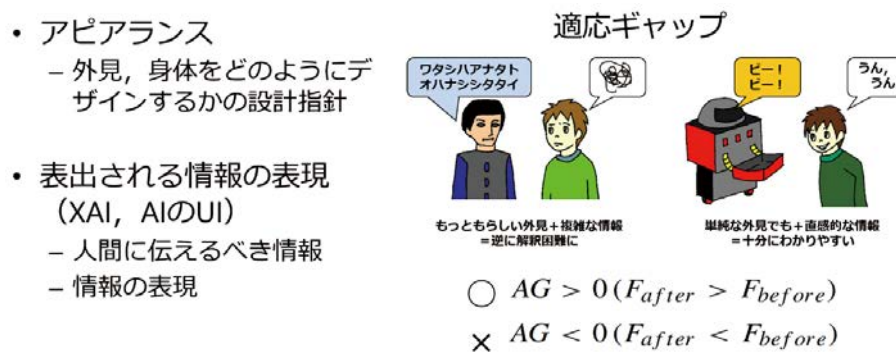


図4-2 エージェント自身のデザイン

エージェント自身のデザイン

エージェント自身のデザイン (図4-2) では、適用ギャップについて考えておく必要がある。外見が人間のように見えるようなロボットに対しては、人間並みの能力を持っているだろうと過剰に期待してしまう。それに対して、素朴なロボット風の外見だったら、対話もろくにできないのではないかと思います。外見から抱くイメージと実際の能力との関係がマイナスだったりプラスだったりするわけで、マイナスの場合は、実際の能力が期待より低いということになるので、デザインとして良くない。実際の能力が外見から推定される能力よりも高いようにデザインする方が望ましい。図4-2に一次不等式を示しているが、HAI デザインでは、モデルを作るだけでなく、このようなフォーマライズも大切だと思っている。

エージェントの機能のデザインでは、どのような機能を持つべきか、機能と外見との関係が適切か、ということに加えて、私が特に興味を持っているのは、どのような学習アルゴリズムをエージェントに持たせるかである。具体的には、学習アルゴリズムの一つとしてMRU (Most Recently Used)、Self-organizing Listとも呼ばれるものがある。これは要するに、一番直近で使った結果を第一候補にするというもので、仮名漢字変換でも使われている。単純なアルゴリズムだけれど、平均計算量的に性能が良いと証明されている。このMRUの良いところは、一般の人でも4-5回インタラクションすれば、どのようなアルゴリズムなのかを、すぐに正しく理解してしまえる、人間に分かりやすい学習アルゴリズムだということだ。つまり、HAI 的には、人

間が適応しやすいエージェントの学習アルゴリズムが重要だと考える。従来の機械学習研究ではほとんどない発想になると思う。

人間とエージェントの間にやり取りされる情報のデザイン

次に、人間とエージェントの間にやり取りされる情報のデザインについて。

説明可能なAI（Explainable AI：XAI）はその一例になるが、インタラクションを非言語情報と言語情報に分けると、HAIは非言語情報に重点を置いた研究が多い。言語情報を扱うことは、自然言語処理や対話システムの研究で行われてきている。なので、顔の表情、音、光、画像、エージェントのアクションなど、これまで取り組みが手薄だった非言語情報の方が、研究としてのブルーオーシャンであるということになる。

人間とエージェントの関係のデザイン

関係のデザインについては、人間とエージェントの間でどのような関係が良いのか、エージェントとの間でもいかに信頼関係を構築するか、そのためにはどういうデザインをすれば良いのかが研究されている。

とても幅があるが、例えば、非常に信頼できそうな外見をデザインするとか、何らかの協調作業をやっている途中で信頼感を感じさせるような振る舞いをするとかいった研究が行われている。シリアスゲーム、ゲーミフィケーション、ヒューマンコンピューテーションなども関係してくる。

HAIの方法論

HAIの方法論は、一般的で統一的な方法はない。その中で、AI・機械学習はコア技術になっていて、少ない訓練例からの学習、能動学習、トランスダクティブ学習、インタラクティブ機械学習などが使われる。特にインタラクティブ機械学習（TAMA、Socially Guided ML）が注目される。

また、HAI研究の基本的方向性として、以下のような点が挙げられる。

1点目はHAI研究で一番よく使われるパターンとして、人間-人間で観察される現象が人間-エージェントでも発生するかを調べるという方法がある。人間X-人間YのうちのYをエージェントに入れ替えるという発想である。

2点目はエージェントがロボットである必然性を疑うというもので、それを問う比較実験が行われる。1点目の方法は、HAI研究だけでなくHRI（Human-Robot Interaction）研究でもよく用いられる方法で、Yをロボットに入れ替えたときに同じ現象が起こるかを観察するというのが行われている。しかし、この2点目の方法は、ロボットありきのHRI研究では見られないパターンである。この2点目は、HRI、マルチエージェントシステム（MAS）、A-lifeの研究との差別化を意識した研究とも言える。

- 信頼関係確立のためのデザイン (trust repair)
- 発達心理学, 社会心理学
 - 心の理論 (TOM: theory of mind)
 - 情動伝染 (emotional contagion)
- 心理学, 哲学
 - Media equation
 - ダニエル C. デネットの3つのスタンス (intentional/design/physical stance)
- 社会科学の説明原理から工学の生成原理へ
 - 表層的な現象を再構成することには疑問

図4-3 最近のHAI研究の動向

最近のHAI研究の動向

最近の動向（図4-3）の一つは、人間とエージェント、あるいは、人間とロボットの間の信頼関係確立のためのデザインをどうするか。特に、一度低下してしまった信頼をどのようにして修復するか、という研究がある。その際にも、やはり人間と人間の間の信頼関係をどう修復するかという研究がベースになっているものが多い。

また、発達心理学や社会心理学、いわゆる心の理論や情動伝染も注目される。心理学や哲学も関わりが深い。既にMedia Equationには触れたが、デネットの3つのスタンス（物理・設計・意図）も有名で、このようなコンセプトを援用するという研究も行われている。

また、HAIに関連する国際会議のリスト^[1]とジャーナルのリスト^[2]を載せておいたので、この分野の動向を知る際の参考にしていきたい。

発表4-2：小野 哲雄

会話システムの発展の歴史

このあと対話エージェントのセッションがあるので、簡単に触れる程度にするが、図4-4に会話システムの発展の歴史を示す。西田先生の会話情報学の資料^[3]から引用した。テキストベースのやり取りから、徐々にリッチなモダリティを用いるエージェントが出てきて、会話システムに発展するという流れが示されている。この図の右側に載せたように、2006年から国内でHAIシンポジウムが始まり、2013年からはHAIの国際会議が始まって、多数の参加者を得ているというのが最近の状況である。

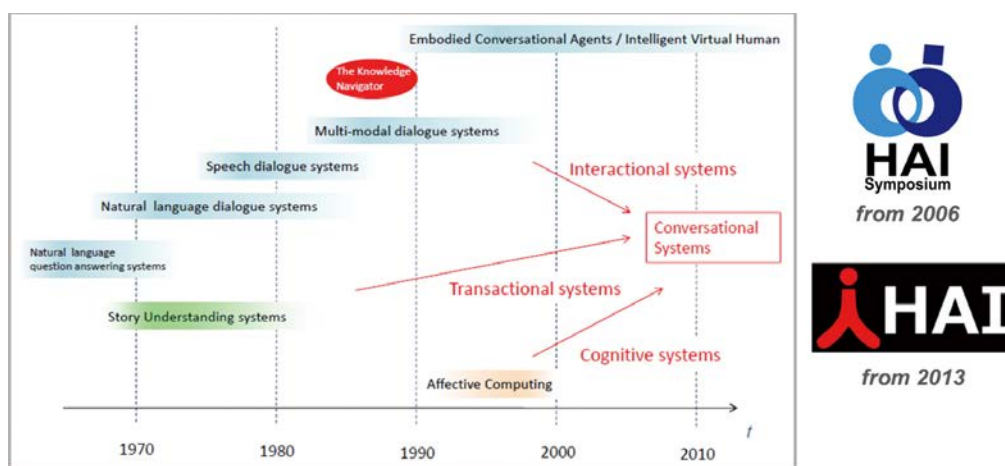


図4-4 会話システムの発展の歴史^[3]

関連する分野やキーワードにも触れておきたい。

まず、環境知能（Ambient Intelligence）が1998年にPhilips Researchから提案された^[4]。これは、人間を取り巻く環境全体を知的にして、人間をサポートするという考えだ。いろいろな分野に波及⁷したけれど、当時はまだ技術的に難しかった。今はIoTや機械学習、そしてHAIなどが発展したので、環境知能は実用レベルになってきている。

次に、Personal AI Agent (PAIA)がIEEEのEthically Aligned Designなどで議論されている。個人データの利用方法とサービスが、データ主体が決めた条件に合致するように、つまり、その個人情報をセキュアにうまく使うと、かなり大きな市場価値を生むというようなことが論じられている。日本でもPAIAの研究が熱心に行われている。

それから、行動経済学^[5]は、限定合理性を踏まえつつ経済学と心理学を足したようなものだ。ここで重要なアイデアとして、ナッジ(Nudge)とブースト(Boost)がある。ナッジはそっと行動を促すというもので、ブーストはよく考えるように促すというものだ。ナッジはよく知られるようになって、Googleのシステムにも取り入れられていたり、日本政府の方針などでナッジを使って効率を高めようとしたりといった話が出ている。

意思決定・行動変容を促すエージェントデザイン

ここから少し具体的に、我々が進めているエージェントデザインについて話したい。

先ほど少し触れたPersonal AI Agent (PAIA)は、いろいろなところで研究されているが、個人情報の管理・運用代行の話が中心になっている。これに対して、我々は、意思決定・行動変容を支援するエージェント（図4-5）を作りたいと考えている。Nudge AgentとBoost Agentを合わせたようなもので、最終的には人のWell-beingを目指そうと考えている。基本的な立場として、リバタリアンパターナリズム、つまり、最終的には人間が判断してほしいという立場でこういうシステムを作っている。

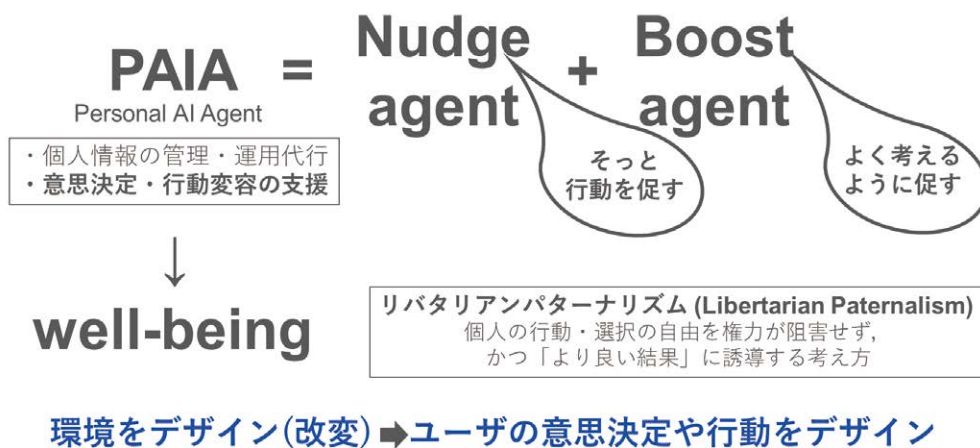


図4-5 意思決定・行動変容を促すエージェントデザイン

先ほど話した環境知能の考え方に基づき、環境をデザイン・改変することで、ユーザーの意思決定や行動をデザインしていこうと考えている。だいぶ前になるが、エージェントマイグレーションに基づくユーザーイン

7 Ambient Intelligenceのコンセプトからつながるものとして、Pervasive Computing、Ubiquitous Computing、Context Awareness、Human-centric Interfaces、Smart Home、Intelligent Roomなどが挙げられる。

ターフェースを用いた環境知能システム「ITACO System」^{[6][7]}を開発した（図4-6）。対話していたエージェントが、ウェアラブルPCからIoTデバイスやロボットなど、いろいろなメディアに乗り移って人間をサポートする⁸。そのとき、エージェントはユーザーの趣味・嗜好を知っていて、例えば照明の明るさなどもユーザーに合わせて調整してくれる。



図4-6 ITACO System

ナッジエージェント（Nudge Agent）

では、我々のPAIAのNudge Agent^{[8][9]}の話に入る。エージェントが環境の中のいろいろなものに移りつつユーザーをサポートする。一つのケースとして、ユーザーが病気のために摂取カロリーの制限をしようと思っているとする。Nudge Agentがいない場合は、例えばトンカツとか高カロリーの食事をつい食べてしまうけれど、Nudge Agentが環境中にいると、さまざまなナッジを使いながら、ユーザーが自分でカロリーの低いもの、例えばサラダや健康食品などを選択するように働きかけるというシナリオで、コンセプトビデオも作った。その中では、低カロリーの食事を取りたいというユーザーの意思を受けて、ユーザーの好みも知ったエージェントが、ロボットやメニューシステムに乗り移り、ロボットから声がけしてお勧めするとか、メニュー表示に視覚的なコントラストを加えて視線を誘導するとかして、ユーザーを低カロリーの食事に誘導しようとする。

ブーストエージェント（Boost Agent）

もう一つのBoost Agentだが、ブーストは、最初の方で述べたように、よく考えるように促すということだ。我々は、天使と悪魔のメタファーを使った意思決定支援のコンセプトビデオ「ADioS : Angel and Devil on the Shoulder」を作った。天使役と悪魔役のウェアラブルロボットが左右の肩に乗っていて、階段を使うかエレベーターを使うかというシチュエーションや、歩きながらスマホを操作するか止めるかというシチュエーションで、天使的なささやきと悪魔的なささやきが左右から聞こえてくる。そのような多様な視点をユーザーに提供することで、ユーザーはよく考えることになり、最終的にはユーザー自身の意思で行動を決めることになる。

図4-7はAna Carabanによるナッジの分類^[10]で、ここではAutomatic – Reflective、Transparent – Non-transparentという2軸で23種類のナッジをマッピングしている。この分類でいうと、我々のNudge Agentは、非常にAutomaticでNon-transparentなタイプのナッジを使って、ユーザーを誘導することを考えている。一方、Boost Agentがそれと逆で、ReflectiveでTransparentなタイプのナッジを使う。この

8 ITACO Systemの動画：
<https://youtu.be/zlWpMMsnDIw>

ような「ナッジ+ブースト」のような考え方が、エージェントデザインの方向性としてあると思っている。

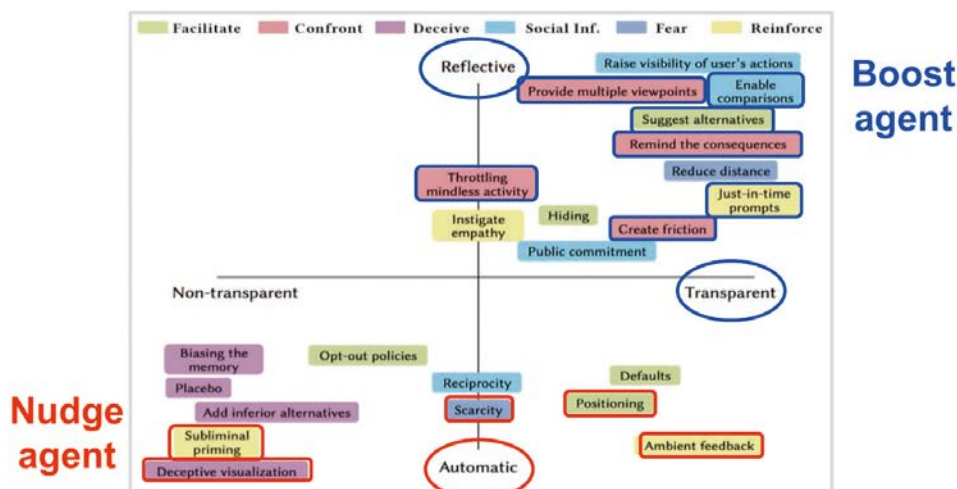


図4-7 ナッジの分類^[10]と Nudge/Boost Agent

討議

大沢 エージェントの信頼性といったときに、その意味には非常に幅がある。人工物を対象にしたときの信頼性として、性能や精度が非常に高いとか安定しているとか、スイッチを切ったら止まるというように自分がコントロールできるとか、いうものがある。しかし、それとは全くの対極のものとして、自分は全くコントロールできないが、その相手の能力・人間性や社会的地位などを認めて、信頼するというのがある。エージェントの自律性を与えるということは、この後者の意味で存在を認めるということになるのではないか、といった議論を自律エージェントの研究の初期によくやったのを覚えている。HAI研究者はこのような点をどう考えているのか。ある意味、社会としてエージェントをどう受け入れていくかという話にもつながってくる。

山田 自律性とは何か、人工物にどのようにして自律性を持たせるか、という問題になると思う。自分としては、自律性を持った存在として信頼するところまでいくのは、まだだいぶ距離があると感じていて、エージェントと人間と同じような信頼関係を築く必要はないと思っている。現状、むしろ前者の方、エージェントに期待する性能と信頼の関数に着目している。その人工物に対する性能ベースの信頼というものがまずあって、その上に、自律性や社会的関係をプラスした信頼関係を考えていくという方向になると思う。

大沢 エージェントは、人間とも違うし、非常に古いタイプのロボットとも違う、何か新しい存在として、人間社会が認めていくことが必要になるのかもしれない。そのようなエージェントが、社会にとって有用なものだと認められた上で。

山田 この話をHAIデザインの立場で考えると、相手の人間がどのような価値観を持っているか、どういう価値観の社会で育ってきたか、といったところまで、インタラクションデザインのスコープを広げるということになる。

小野 コミュニケーションというものを考えると、工学ではシャノンのコミュニケーションモデルといったものがあり、一方、人間と人間の場合は「心の理論」(Theory of Mind)があり、かなり複雑なプロセス

として捉えている。HAIはまだ統一された方法論がないのだけれど、それらの中間ぐらいだという感覚があり、人間のインタラクションについて、そのときの人間の思考がどうかとか、どういうプロセスで行動するのかとかも考えようとしている。これからの社会で、人間と人工物、特にエージェンシーを持った人工物とやり取りする機会は絶対に増えていくと思うので、今議論になっているような面を考えていくことは必要だと思う。

大沢 そのときに、理解するということと物を作るということでは、やはり社会的責任が異なると思う。理解することには責任がないというわけではないが、物を作って社会に放つと、それがもたらすプラスの価値、経済的な効果の一方で、リスクや損失もある。だから、HAIも工学の一分野として物を作っていく以上、社会に対するインパクトは必ず考えておく必要があるし、エージェントの信頼も含めた人間社会との関わりをどう理解していくのかは結構重要な問題だと思う。

福島 デジタル社会における信頼（トラスト）の問題については、ちょうど我々のところで俯瞰的な整理を試みた報告書^[11]をまとめたので、参考にしていただけるかもしれないし、ぜひ別途議論もお願いしたい。

野田 RoboCupで人間とロボットが相対するときに、ロボットの意図が読めなくて非常に困ると言われる。それはロボットに関して、表情・目線や身体の重心移動のような情報を人間がつかめなくて、ロボットの次の動きが予想できないからだ。現状はそうだけれど、ロボットには物理的な身体があるので、それをうまく使うことで改善できると期待している。しかし、仮想エージェントは物理的な存在ではないので、相手に働きかけるアクチュエーターになるものは、やはり言葉が中心なのではないかと思う。仮想エージェントの外見として、目や表情を使うことは可能だが、ロボットを二次元にただけということになる。そうではないもので、エージェント独自の働きかけとして、言語以外のものが何か考えられるか。

山田 それはまさにHAIの表現デザインの問題として、いろいろな取り組みがある。ロボットの場合は今の話にあった通り、動きが重要で、サッカーのフェイントとか、人間は相手のちょっとした動きからどう動くかを予測するというものを使っているわけだ。なので、エージェントがそういった動きの一部を意図的に見せることで、例えば、自分のゴールがどこにあるのかとか、複数あるゴールの中のどれを自分が今ゴールとしているのか、といったことを伝えようという研究がいくつかある。そういったことを込みで協調プランニングをするとか、ベイズ推定を使うとか。動き以外では、私のグループでやったものだが、ビープ音の周波数変化で確信度の上下を伝えるとか、他にも、光の明滅パターンでいくつかの内部状態を伝えるとかいったASE (Artificial Subtle Expression) と呼ばれる研究事例がある。ただ、言語ほど伝達するコンテンツをリッチにはできなくて、ほとんどの場合は1/0のバイナリー情報のように解像度が低い。しかし、解像度が低いとしても言語に依存せず、ある種口バスタな伝わり方なのは良い点だ。

野田 もう一つ質問は、人工物のトラストに関することで、例えば、自分と自分の子供では、スマホとかスマホのアプリに対する信頼度が全然違う。自分はまだ古い人間なので、大きいコンピューターじゃないと安心できないとか、キーボードじゃないと安心して入力できないといったことがあるが、子供たちはもうスマホが当たり前で、すごく信頼している。そのような周りの環境の変化によって、人間の信頼が変化してしまうことについて、HAI研究ではどう考えているのか。

山田 トラストが、その時点でのテクノロジーのレベルや、各人がテクノロジーと接触している経験など、社会的な背景や経験に依存するというのは、その通りだと思う。だから、例えば、擬人化エージェントや仮想エージェントに対してどれくらい信頼を置くかとか、スマートフォンに対してどれくらい信頼を置くかというのは、いつどのような参加者で実験をやるかによって変わってくる。なので、例えば、それまでロボットを触ったことがあるか、家でルンバを使った経験があるかというようなプロファイリングをした上で、こういう結果が出たというような研究のやり方になる。

5 | 対話エージェント領域

発表5-1：東中 竜一郎

最先端のテキスト対話例

対話エージェント領域は、私がテキスト対話システムを中心に話し、駒谷先生がマルチモーダル対話を中心に話すという分担にした。

最初に、SimpleTODという最近の対話システムによるホテル予約対話の例を示す（図5-1）。このようなシステムは20年くらい前にもあったが、ホテル、レストランなどマルチドメインで会話ができるという点と、全部ニューラルネットで動いている⁹という点が異なる。

User 中心部のホテルを探しています。
 System 中心部にはホテルが3つあります。価格帯はいかがなさいますか？
 User 4つ星のホテルを探しています。
 System 申し訳ありませんが、ご希望の条件に合うホテルは中心部にはありません。
 User 3つ星ではどうでしょうか。
 System 1件見つかりました。予約しましょうか？
 User ホテルの名前を教えてください。
 System Gonville プレースにある Gonville ホテルです。
 User 予約をお願いします。
 System 承知しました。何人で泊まりますか、また、何日から何日までの滞在を希望しますか？
 User 2人で2泊します。
 System 何日に到着されますか？
 User 水曜日です。
 System 予約は成功しました。参照番号は76896です。
 User ありがとうございます。さようなら。
 System ご利用ありがとうございました。

図5-1 タスク指向型対話システムとの対話例¹⁰

対話例をもう一つ示す（図5-2）。私が主催している「対話システムライブコンペティション」で、昨年最も良い成績だったシステムによる対話例で、いろいろな話題に追従した会話ができている。このシステムはLINEが開発したHyperCLOVA（GPT-3のクローンのようなシステム）をベースに動いている。英語圏で評判の良いBlenderBot2という、Facebookが開発したシステムでは、インターネット検索をしながら対話をするようになっている。このようなシステムが今の最先端である。

⁹ SimpleTODはGPT-2を用いている。

¹⁰ SimpleTODとの対話は英語だが、それをDeepLで英日翻訳したものを載せた。

System 初めまして、私はよつばです！さて、何のお話をしましょうか。

User よろしくお願ひします。

System こちらこそよろしくお願ひします！趣味はあるんですか？

User そうですね。車など好きです。

System 車良いですね。ドライブ楽しいですもんね。ちなみに車種は何に乗っているんですか？

User 今は違いますが、ずっとホンダ車に乗ってました。

System なるほど。確かにホンダ車は丈夫だと聞きますね。これからも大切にしてくださいね。

User はい。ところで本田宗一郎をご存知でしょうか。

System もちろん知ってますよ。彼がいなければ自動車産業はここまで発展しなかったと思います。

User そうですね。偉大な方ですね。

System 本当にその通りです。世界中の人々に夢を与えてきた素晴らしい人物です。

User ええ、ひと昔前のF1などではホンダエンジンが天下を取ってましたものね。

System あの頃は面白かったですね。セナ・プロスト時代でしたっけ。

User よくご存知ですね。マクラーレンホンダ有名ですね。

System ああ、懐かしいですね。当時小学生だった私でも興奮した記憶があります。

User 同じくです。セナは惜しいことをしました。

System はい。残念ながら事故で亡くなってしまいましたからね。

User 話は変わりますが、大相撲は興味ありますか。

System もちろんですとも！千代の富士全盛期世代ですから。ちなみに好きな力士は誰ですか？

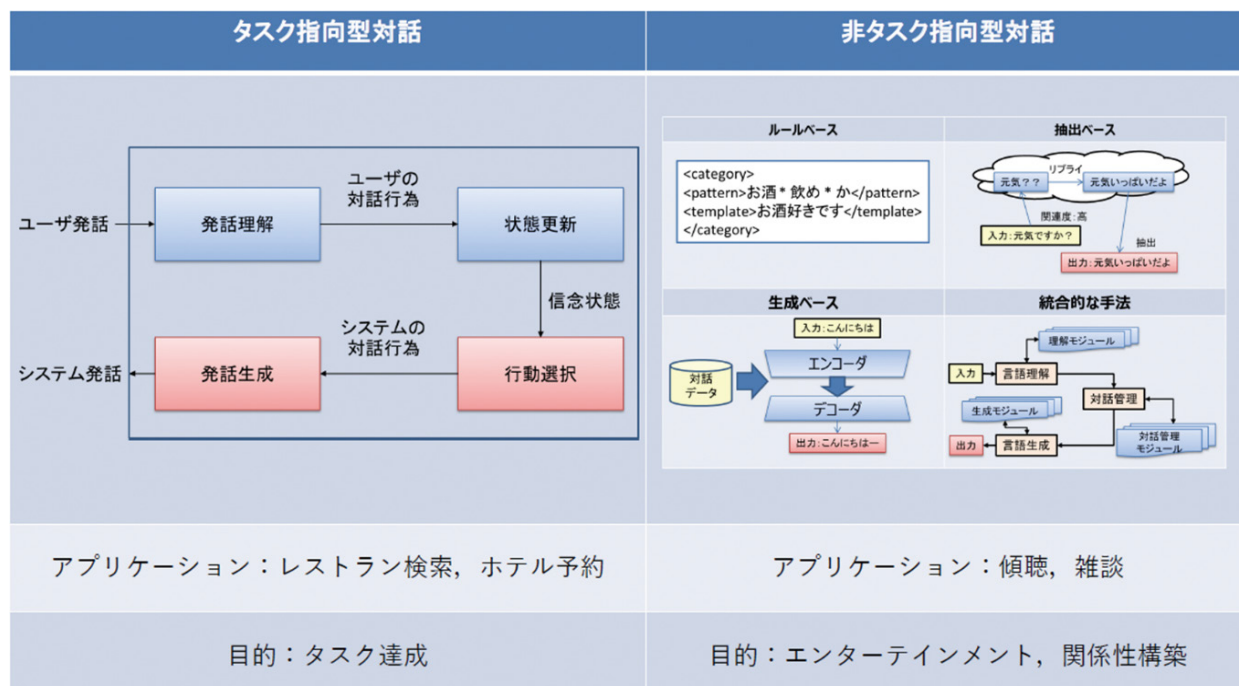
図5-2 非タスク指向型対話システムとの対話例^[1]

図5-3 対話システムのアーキテクチャー

対話システムのアーキテクチャー

ホテル予約の例のようなタスク指向型対話を行うシステムのアーキテクチャーは、昔からあまり変わっていない。図5-3の左半分に示すように、ユーザー発話が入ると、それを理解して、状態を更新して、何をするかを決めて、何を言うかを決めたら、それを言語化するというものである。ただ、これらのコンポーネントや

システム全体が深層学習・ニューラルネットに置き換えられている。アプリケーションとしてはレストラン検索やホテル予約のようなものが多い。それらのタスクの達成率を上げることが目的になる。

これと比べると、非タスク指向型対話の方は、まだ方法論が固まっていない。図5-3の右半分に示すように、いろいろな手法がある。ルールで書くのがルールベースの手法、検索してその結果を使うのが抽出ベースの手法（検索ベースともいわれる）、完全にニューラルネットで行うのが生成ベースの手法、それらを組み合わせるのが統合的な手法である。生成ベースを言語生成に用いつつ、最終的に何を言うかを抽出ベースの手法でランキングするというのが、実用的なシステムでよく使われている。その中にはルールベースの手法も結構入っている。非タスク指向型のアプリケーションとしては傾聴や雑談、癒しとか、時間をもたせるときとかに使われている。エンターテインメントや関係性構築などが目的とされる。タスク指向型対話であっても、最初の方に雑談を入れると、全体の効用がよくなるとか、信頼感が上がるとかということも知られている。

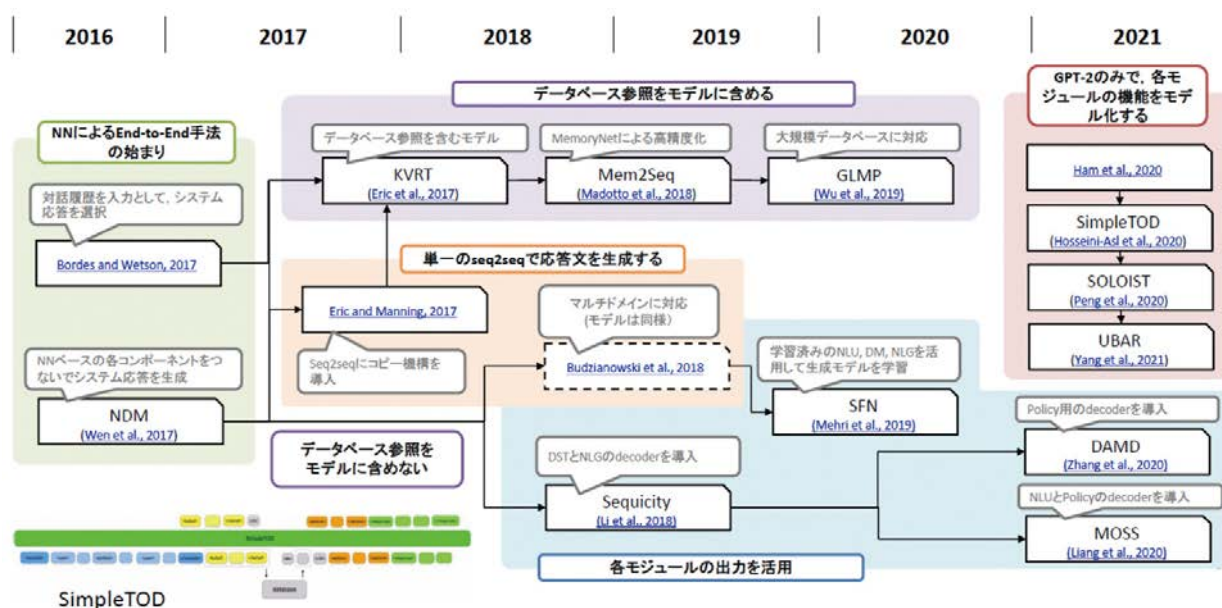


図5-4 タスク指向型対話の系譜

タスク指向型対話

図5-4にタスク指向型対話の系譜として、ここ5-6年の流れを示した。2016年頃からタスク指向型対話でニューラルネットが使われ始めた。タスク指向型ではデータベースを参照して、その結果をもとに生成するが、そのデータベース参照も含めて全部1個のニューラルネットモデルに含めるというのが、図5-4の上側の流れになる。

それに対して、それ以外のところで頑張るというのが、図5-4の下側の流れである。そこでは、各モジュールの出力を活用するというのがポイントになっている。End-to-End全てニューラルネットでやろうというのはやはり無理で、内部的には信念状態やシステムのアクションなどをつなげていくことが必要だと考えて、モジュールをニューラルネットで置き換えつつも、それらを統合して、全体で良い性能を出すようにしてきた。

2021年に出てきたのが、1つの言語モデルGPT-2上で全部をやっってしまうというものである。どういう信念状態を出せばよいかをまずアウトプットさせて、それを使って外部のデータベースを検索して、それをまた入れ込んで、そこから次に何をやるか決めて、何をやるか決めたら、それに関してシステム発話を生成するといった一連のことを、GPT-2上でやっってしまうことができています。しかし、このEnd-to-Endの枠組みでは、まだ70%程度のタスク達成率しか得られていない。人手で作れば90%以上になるので、まだまだである。

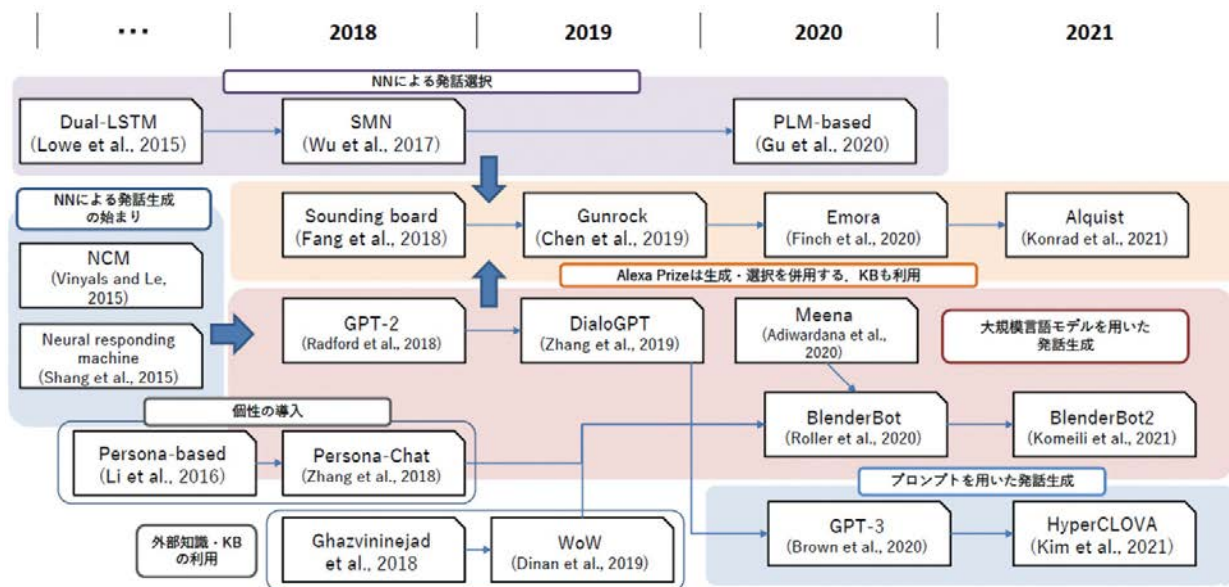


図5-5 非タスク指向型対話の系譜

課題としては、状態更新の精度改善、発話生成の精度改善、End-to-End全体でどうモデル化するか、が挙げられる。状態更新では複雑な情報を扱うための学習の効率化、発話生成ではハルシネーションの抑制が課題である。ハルシネーションというのは、入力にない情報をニューラルネットが勝手に言ってしまうという問題で、どうすればこれが抑制できるかが盛んに研究されている。

非タスク指向型対話

一方、非タスク指向型対話の系譜を図5-5に示した。中央のラインはAlexa Prizeというコンテストで、どのようなシステムが作られてきたかを示している。生成モジュールと選択モジュールを併用し、かつ、中身にはナレッジベースもどんどん取り入れていくという方式が取られていて、この方式はここ4年くらい変わっていない。その下のあたりは、大規模言語モデルを使った対話モデルである。大規模な対話データを使って発話生成をするもので、BlenderBotやBlenderBot2が有名である。

さらに最近では、図5-5の右下あたりにあるGPT-3やHyperCLOVAのようにプロンプトを使った発話生成が注目されている。ほとんど学習せず、例えば冒頭部分だけを書くと、残りが勝手に生成される。こういうときにはこういうふうに話すという機構を使って対話ができ、さまざまなドメインに適用しやすいことから、このような機構を使った研究が少しずつ出てきている。

非タスク指向型対話システムは、既に15ターン程度であれば人間かどうか分からないレベルのチャットができている。しかし、個性・論理の一貫性や、外部知識の利用は課題である。15ターンを過ぎていくと少しずつほころびが出てきて、2回目や3回目の会話になると、また初対面のときのような会話になってしまうという問題点がある。外部知識の利用では、最新の情報にうまく答えるようにするとか、モデルに含まれない内容に回答できるようにするとかが課題として取り組まれている。

その他にも、読解や画像処理と融合した対話システムを作ることや、システム評価の仕方も課題である。システムが学習した言語モデルに基づいて発話した内容が倫理的な問題になってしまうとか、ステレオタイプのことを言って問題になってしまうといったことも起きている。

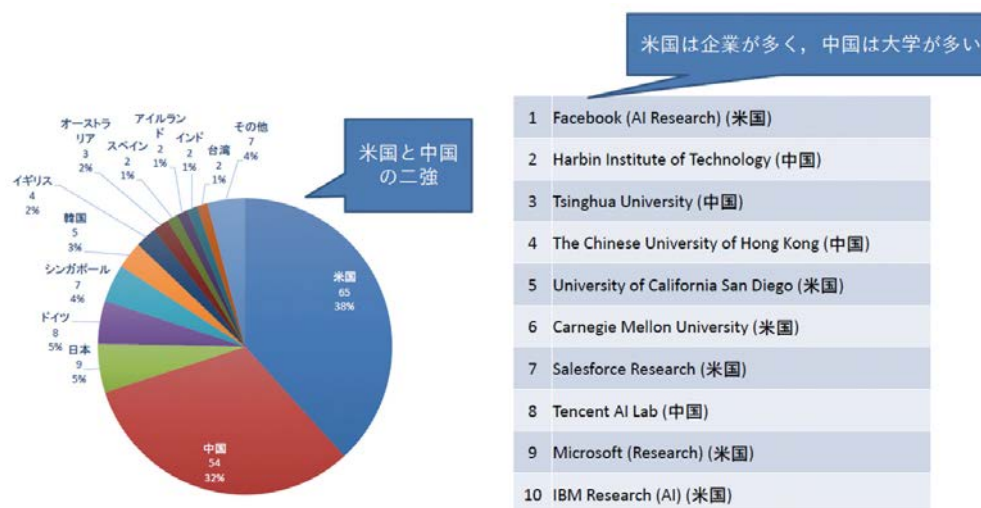


図5-6 国別の動向

最近の動向と今後の展望

コンペティションからも最近の動向^[2]が分かる。Alexa Prizeも最近ではチャットからタスクに移行してきている（TaskBot）。実世界のタスクでアシスタントができるようなものにもチャレンジしている（SimBot）。あとはDSTC（Dialog System Technology Challenge）、ある特定のシチュエーションでうまく対話できるとか、マルチモーダル情報を使って対話するといったコンペティションが出てきている。テキストからマルチモーダルへ徐々に移行しつつある。

図5-6には国別の動向^[3]を示した。やはり米国と中国の二強になっている。米国の企業が非常に元気で、中国の大学も頑張っている。

最後に今後の展望だが、タスク指向型対話については、音声コマンドからスロットフィリングへ発展し、現状はまだ特定の話題で話すということしかできていない。さらに定型タスク、そして不定形タスクへの発展が今後の方向性である。非タスク指向型対話については、一問一答から話題の共有というレベルへ発展してきたが、さらに情報の共有、そして価値観の共有への発展が今後の方向性である。窓口対応のようなうまく情報状態が定義できないようなタスクや、実際の何らかのエピソードなどを共有するといったタスクはまだできていない。なぜできないかというと、お互いの理解を積み重ねていくような、いわゆる長期に対応するためのデータ構造・理解構造といったものがないからだ。ユーザーとシステムの共同作業の研究が増えてきているし、新学術領域の「対話知能学」では概念のすり合わせに着目している。そういう共通理解（共通基盤とも呼ぶ）の研究というものが今後重要になると私は考えて研究を進めている。

発表5-2：駒谷 和範

マルチモーダル対話システム

今、世の中に広まっている対話システムとして、スマホ上のアプリ、Amazon EchoやGoogle Homeのようなスマートスピーカー、Pepperのような対話ロボットなどがある。これらは音声認識結果のテキストのみを使っている。

これに対して、「テキスト+別のモダリティー」を使用した対話システムを、マルチモーダル対話システムと

呼ぶ。対話システムなので言語情報を基本にしつつ、別のモダリティー（複数あってもよい）をプラスする。「テキスト+画像」が多いが、さまざまなセンサー情報を使うこともある。例えば、音声を使う場合でも、韻律情報、つまり声の様子から感情を推定するという研究がある。センサー情報として筋電、触覚、生理指標などもあり、人間の視聴覚に限る必要もない。

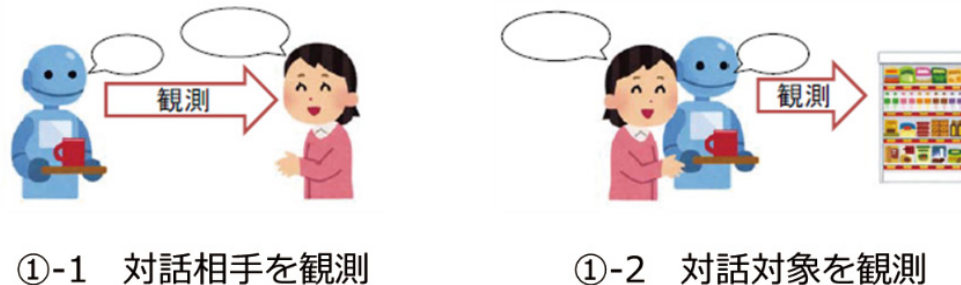


図5-7 状況認識の種類

マルチモーダル対話システムの研究課題は、①入力に関する状況認識と、②出力に関する表出に分けられる。入力については、観測対象によって、対話相手を観測する場合（①-1）と、対話対象を観測する場合（①-2）という2種類に分けられる。ロボットと人が対話しているとき、人をマルチモーダルで観測して対話するというのが①-1で、対象と一緒に見ながら、その対象についてあれこれ対話するというのが①-2である（図5-7）。

①-1 対話相手を観測

まず、対話相手の観測については、さまざまなセンサー情報から人間の内面状態を読み取る、社会的信号処理（Social Signal Processing）と呼ばれる分野がある。人間の内面状態としては、端的には感情、他にも態度、リーダーシップ、参加度（エンゲージメント）などが考えられる。必ずしも対話に限らないが、社会的信号処理の研究やコンペティションが、ICMI（International Conference on Multimodal Interaction）などで行われている。

マルチモーダルの中でも特に音声については、音声からパラ言語を認識するということが行われている。言語に含まれる文字以外の情報をパラ言語という。Interspeechという国際会議で、Computational Paralinguistics Challenges（ComParE）という特別セッションが2009年から毎年開催されている。音声の会議なので入力は音声だが、対象はさまざまで、2009年は感情、2010年は年齢や性別、2011年は眠たさだった。2021年はコロナの咳の音を判別するというチャレンジだったらしい。

アプリケーションとしては、カウンセリング対話や面接の練習が挙げられる。2013年に南カリフォルニア大学で開発されたSimsenseiというシステムのデモ動画をYouTubeで見ることができる¹¹。この動画では、左側にユーザー、右側にカウンセリングをするエージェントが表示される。言語だけでなく、ユーザーの表情、声色、話すタイミングなども見ることで、この人は少し落ち込んでそうだとか、この話題は楽しそうだとかを判別しながら、エージェントがユーザーからうまく話を聞き出すというようなシステムである。

11 “SimSensei & MultiSense: Virtual Human and Multimodal Perception for Healthcare Support”
<https://youtu.be/ejczMs6b1Q4>

①-2 対話対象を観測

次に、対話対象の観測の例の一つとして、SIMMC (Situating and Interactive Multimodal Conversations) のデータが公開されている。これはタスク指向で、買い物をタスクとしている。DSTC (Dialog System Technology Challenge) というコンペディションの中の1トラックになっている。例えば、図5-8のような画像を見ながら、エージェント (システム) とユーザーが買い物について話すことが想定されている。このデータには、このような商品画像に加えて仮想空間の映像も含まれている。このとき、画像に対する照応解析が必要になる。例えば、このジャケットについて話していて、「その下にあるスカート」と言ったら、それが何を指すか分かる必要がある。

Situated Context	Dialog Utterances	Dialog Annotation
	U: What can you tell me about the style of <u>that brown skirt</u> ?	[IN:ASK:GET:SKIRT.skirtStyle What can you tell me about the style of [USER.attentionOn this] skirt?]
	A: This style is loose ball gown.	[IN:INFORM:GET:SKIRT.skirtStyle This style is [O.skirtStyle loose ball gown]]
	U: Can you show me another brown skirt?	[IN:REQUEST:GET:SKIRT Can you show me [O.sequential another] [O.color brown]skirt?]

図5-8 SIMMCデータの例^[4]

応用シーンとしては、メタバース内での対話、RoboCup@Homeのようなロボットへの指示、自動運転車への指示などが考えられる。自動運転車に「次の信号を左に曲がって」と言って、それをきちんと理解してくれると非常にうれしいが、これは格段に難しい。図5-8のケースは固定画像が対象だが、自動運転のケースでは、カメラから入ってくる環境の映像が動く中で高い状況認識精度が求められる上に、車の位置を踏まえたタイミングでの処理が必要になるので難しい。

②表出

次に、表出については、実環境で対話の相手となる物体として、ロボットや仮想エージェントなどがある。既に話のあったHAI (第4章) で多くの検討があり、知見も蓄えられている。

ここで紹介したいものとして、表情で感情を表出するロボットがある。スウェーデンのFurhat Roboticsという会社が開発した首だけのロボットは、顔に裏からプロジェクションマッピングすることで豊かな表情を出せる¹²。日本にも代理店があり、バンダイナムコやJR東日本とのコラボレーションも行われている。

マルチモーダル対話システムの技術課題

技術課題としては、まず、マルチモーダルのシステムを作ることの難しさが挙げられる。入力となるセンサーが複数になり、個人差や環境変動もあるので、実環境で頑健性を確保することがとても大変になる。さらに、

12 Furhat Robotics 社の Web サイト
<https://furhatrobotics.com/> でその外観などが見られる。また、バンダイナムコやJR東日本とのコラボレーション事例については、こちらの記事
<https://jp.techcrunch.com/2019/08/20/bandai-namco-research-furhat/> で紹介されている。

対話は基本的にリアルタイム処理が必要で、相手や対象が動くとなますます難しくなる。ターンテイキングでは、少し話すのが遅れるだけで、そこに含みが生まれてしまうという問題もある。音声と言語と画像だけでも、それぞれのフレーム単位も異なるので、それらをどのようにして統合するか、複数モダリティー間の統合にも難しさがある。

もう一つの技術課題として評価規準が挙げられる。テキスト対話システムでも課題だが、パラ言語情報が入ってくるので、さらに難しくなる。例えば、最近は感情認識用対話コーパスとして皆が使っているIEMOCAP (The Interactive Emotional Dyadic Motion Capture)¹³は、役者がその感情を演じたデータである。なので、役者が怒っているように演技したら、「怒っている」と判定するのが正解となる。しかし、自然な対話では、いきなり極端に怒るというものではないし、楽しそうにしている場合でも、第三者が複数人で評定すると必ずしもその結果は一致しない。正解データが一意に定まらないので、機械学習にも乗せにくい。そもそも、何がよい対話なのか、対話自体の評価も難しい。

ということで、マルチモーダル対話システムは、いくつかの優れたデモシステムが存在し、機械学習を用いた方法論も導入されつつあるが、ここで挙げたような難しさがあり、統一的な枠組みの開発はまだ始まったばかりという段階である。

討議

大沢 最近の機械学習ベースの自然言語処理の技術発展は著しいが、意味を理解するということを諦めているのではないかと、そして、意味を理解していないシステムを人が信じることで、何か本質的な問題が起きてしまわないかと危惧している。

東中 意味は理解したいと思っているし、タスク指向型対話は、内部的な意味フレームのようなものにグラウンディングしているからこそ、いろいろな対話状態に応じた対話ができるのだと思う。

大沢 タスク指向型はたぶんそうだと思うが、自由会話ではどうか。例えば比喩を使うときなど、個人の体験に基づいた意味構造や関係性に依存すると思うが、内部モデルの系が爆発しそうで、かなり難しいことのように思える。そういったところは、最近どう扱われているのか。

東中 今はとにかく大規模にやればいいという傾向にある。しかし、先ほど話したように、共通理解が重要になると考えている。相手のモデルを持ちながら対話するとか、マインクラフトなどで共同作業をするといった研究が出てきている。ニューラルネットが内部モデルに使われているが、相手が何を考えているかは、逐次アンケートを取っているというレベルだ。相手が今何をしようとしているとか、相手がどんな知識を持っているとかかを推定しながら動くようにするというのが、今の共同作業系の研究で一番頑張っているところだと思う。しかし、ハルシネーション問題が起きてしまうのも、内部で意味的にグラウンディングできていないため、現状は、検索した結果と突合することで、なるべくおかしいことを言わないようにしているというレベルだ。

大沢 それはシステム自体に説明能力を持たせるという話につながる。そのため、内部にできたものが何なのか、人間が理解できるようにするなど、機械学習の中にいろいろなものを組み込んでいくという取り組みや、そういったことを包括した研究が、これから必要だと思う。

東中 タスク指向型対話と雑談（非タスク指向型対話）はかなり融合してきていて、雑談でも、やはりしっかりとしたモデルを内部に持って、そこにグラウンドしながらの対話という問題が、徐々に扱われていくだ

13 <https://sail.usc.edu/iemocap/>

ろうと思っている。説明能力についても、人間も完全に説明できるわけではないが、必要な一定のレベルのようなものはあると思う。

山田 ニューラルネット End-to-End の対話システムでは、性能がそれほど出ないということだったが、画像認識話ですごい性能が出るのと何が異なるのか。やはりある程度トップダウンにモジュール構造のようなものを与えないと無理だということか。

東中 現状一番強いのがパイプラインで、理解系だけ Bert を使って、他は全部ルールで書かれたシステムだと 95% が達成できる。それに対して End-to-End では 70% しか出ない。人間がしないような勘違いを結構してしまう。タスクをコンプリートするには、ある種の決定的なところが必要だが、そこを扱っていない。一方、End-to-End も徐々に性能が上がってきているが、データ作成がものすごく大変だという問題もある。MultiWOZ がほとんど唯一のデータセットで、皆これで数年ずっと研究している。信念状態を全部ラベリングして、ユーザーごとの発話を全部ラベリングした上で、各コンポーネントのロスと、全体のタスクのロスを最適化するように動かす。そのデータがないとできないのだが、まずクラウドソーシングを使った後、さらに人間がチェック・修正して作っている。

大沢 画像認識との違いについて言うと、画像認識は対象物が確定しているのに対して、言語理解や言語によるコミュニケーションは、実世界にある対象物を扱うケースもあるが、大半は共有信念の形成なので、ずっと難しいのだと思う。

福島 マルチモーダル対話研究から HAI 研究を見るとどうなのか、逆に、HAI 研究からマルチモーダル対話研究を見るとどうなのか、異なる面が強いのか、相互にどう絡んでくるのか。

駒谷 2 つはとても関連している。自分の発表でマルチモーダル対話の評価の問題を挙げたが、人間でも、口では楽しそうに話しているけれど表情がとても硬いとか、何か話し方の印象が悪いとかいったことがあるように、対話システム全体の評価では、話の内容に加えて、見た目や振る舞いによっても印象は大きく左右される。そういう意味でも HAI の知見は非常に大事だと思っている。

山田 いわゆるパラ言語が関連して、例えば、はっきり明瞭に「分かりました」と言うのと、弱くぼそっと「分かりました」と言うのとでは、本当に分かっているか、実は分かっていないかというくらい全然意味が違う。そのようなテキストだけでは意味が決まらないというシチュエーションはよくある。なので、マルチモーダルでやるというのは適切な方向だと思う。HAI 自体は基本的にマルチモーダルなので、駒谷先生の研究とは非常に近く、半分以上かぶっているように感じる。

6 | 総合討議

マルチエージェントシステムと社会

福島(司会) AIの今後の発展として、一つのAIの中身をどう作るだけでなく、AIとAIの間のインタラクションや、AIと人のインタラクションを考えることや、それによって社会の中の活動がどう形成されていくかを見ていくことは、ますます重要になると思う。このような今後のAIへの期待に対して、今日のエージェントに関わる3つの研究領域は、一見だいぶ異なるように見えても、実は非常に深く関係してくるものだと思う。3領域の融合や社会との関係についてどのような感触なのかを聞いてみたい。

大沢 その方向性に関して思ったことだが、マルチエージェントで社会メカニズムをどう作っていくかという話があり、マルチエージェントによる社会シミュレーションの話もあり、そういった研究と経済モデルとの関係性はどうか。それらは扱っているものが似ていると思う。

伊藤 今までのマルチエージェントの研究は、ゲーム理論やメカニズムデザイン理論に寄り過ぎていたと思う。基本的にはフォン・ノイマン・モルゲンシュタイン効用、プロスペクト理論などを経て、行動経済学の結果を取り入れた効用理論ができてきていて、理論的には人間の効用の数学的なモデル化ができつつある。その中で、従来のゲーム理論でいう超スーパー合理的プレーヤーだけではなく、ソーシャルネットワークの中での価値観や他人の価値観など、もう少し他人のことを考えながら自分の効用を最大化していく、社会の効用を最大化していくというモデルも徐々に受け入れられつつあると思う。そうすると、今までの社会経済学やメカニズムデザインで言われている社会システムとは違ったものができ得ると思っている。それは人間同士の協調というところにつながってくる。その効用関数や社会の仕組みのようなものを、どのようにして現実の日本の人々と共有していくか大事になってきて、自然言語処理技術を高めていく必要もあると感じている。

福島 経済学の方でも、以前は、人間は合理的な行動をするというモデルを作っていたが、やはりそれでは限界があるので、個人の揺らぎや他者への配慮など、人間的な側面をいろいろ入れていかないと、経済のモデルとしても不十分だという認識が出てきていると聞いた。

大沢 今回のような機会から今後の提言につなげることを意識すると、社会システムやそのメカニズムの設計に、経済科学やAIがこんなふうに関与できると示唆できる可能性があると思う。

伊藤 それはまさに私が言いたいと思っていることだ。

野田 その思いはよく分かっているし、共有しているところだ。実際、自分は交通系でいろいろやっているが、既存のサービスとぶつかるときになかなか理解されない、抵抗勢力の壁をなかなか破れないという問題にぶつかる。COVID-19でもAI研究者は一生懸命シミュレーション分析をやって、Gotoキャンペーンもオリンピックもかなり危険だという分析を出していたのに、それが具体的な判断につながっていなかったという問題を、どうやって改善していったらいいのか。いろいろと証拠・事例を見せるしかない、小さなところからでも成功事例を積み重ねるのが一番説得力を持つというのが現在の私のスタンスだ。

マルチエージェントとHAI

福島 3領域の融合については、特にマルチエージェントの側からすると、HAIのような取り組みはだいぶギャップがあるのか。将来的な方向感としてはつながりそうに思うが、今日話を聞いたところではまだこれからだという印象を持った。何が課題なのか。

伊藤 HAIで取り組まれているようなことを、マルチエージェントの側でもやりたかったのだと思う。しかし、研究分野としては、数学的なモデル・背骨のあるところから進むのだと思う。課題としては、トップカ

ンファレンスにトレンドがあって、そのトレンドに乗っていない研究がなかなか通りにくい。今だとAI全般は深層学習や機械学習に寄っていて、マルチエージェントのトップカンファレンスはゲーム理論やメカニズムデザインにだいぶ寄っていた。ただ、今からその反動が来るだろうと私は思っている。そこで、私はHAIのビジョンが重要になると思う。

福島 そこがつながってくるのであれば、日本発の研究分野だということが強みになる。

大沢 HAIや対話は個別のインタラクションでマイクロスコピックだが、マルチエージェントも初期はマイクロスコピックなところから始まったものの、今はたぶんメゾスコピックな規模になっている。そういった規模によって方法論が変わってくるから、メンタリティーも違っていると思う。

山田 マルチエージェントとHAIの関係について、少し引いて見てみると、伊藤先生の研究は、人間と人間のインタラクションをエージェントがメディエートする研究のように思える。つまり、私の発表で、HAIでは3種類のインタラクションを対象にしようというのの一つに該当する。それで、そのメカニズムデザインは、インタラクションをデザインしていると解釈ができると思う。それから、AAMASやNeurIPSでも、インタラクティブ機械学習とか、人間が教師役になってエージェントの学習を導いていくという研究は、必ずしもメジャーではないけれど脈々とある。私がHAIを始めたときは、人間が系に入っているマルチエージェントシステムの研究は皆無だったが、現在は変わってきているのかなと思う。そういう意味でも、今回の集まりはありがたかった。

伊藤 マルチエージェントシステムのコミュニティとしても、10年以上前はエージェントに人間が入っていなかったけれど、2010年前後にエージェントに人間も入れようという話をしたので、今はもうそのように考えていると思う。

野田 マルチエージェントシステムを動かしていると、いいインターフェースが欲しいと思う。例えば乗り合いのケースで、調整はこれから結構大事な技術になると思っていて、既に配車を決めた客に対して、少しずらしてもらえませんかと言いたいときなどに欲しくなる。それを入れたプランニングも研究したいが、そもそもどんな出し方があるのかのオプションが分からないから、その問題を定式化できないでいる。何か面白い道具立てが提供されるとうれしい。マルチエージェントで何かやるというのは、最終的に人々に対して何かを促すことになる。その一番小さいループは個人に対するインタラクション、まさしくHAIになると思う。

森永 産業的には、HAIが使い物になるのであれば、自動交渉のターミナルのところに絶対入れることになると思う。それから、時期については、メタバース的な世界が広まるなら、結構すぐ来ると思うけど、そうでなかったら、人に対するターミナルを普及させるところが重くて、産業的には厳しい気がする。

エージェントの振る舞いの制御

福島 メカニズムデザインの的にエージェントに対して外から社会のルールを与えてうまく動かすアプローチと、エージェント自身が何らかの効用関数に基づいて振る舞うようにするというアプローチは、今どういう関係になっているのか。メカニズムデザインの的なものは、ある分野ではうまく回っているけど、何か限界があるような気もしている。個々のエージェントがどんな意思を持って振る舞うかをうまく盛り込まないと、いろいろな社会活動・現象をカバーできないのではないかと、何となく直感的に思った。そのあたり、学術的な発展がありそうか。

伊藤 基本的にはゲーム理論で出てくる数字は全て効用関数で、フォン・ノイマン・モルゲンシュタイン効用を仮定していて、それは所与のものだということになる。その応用においては、例えばポーカーであれば、ポーカーの手に点数が付く。セキュリティゲームの警備では、実際にどれぐらいの価値があるか、インタビューで聞き取り、ゲームマトリックスを作っている。その意味で、効用関数をどのようにすべきかについては、経済学の中でかなり議論がある。エージェントでは、BDIアーキテクチャーとか、コグニティブアーキテクチャーとかやられていたので、そのあたりの知見も取り込めると思うが、マルチエー

ジェントとしての挙動を数学的モデルでどう担保していくかは、学術的に問題があるかもしれない。

福島 対話エージェントから見て、この話は何かつながるところはあるか。

東中 最近の学会でちらほらと見るようなキーワードが目につく。例えば信頼であるとか、複数人の対話ユニット間での対話とか、交渉であるとか、割と新しいタスクとして提案されてきているので、接点は近づいているという印象を持った。それから、タスク対話に取り組んでいる我々としては、例えばユーザーシミュレーターとか、シミュレーション系の技術がもう少し使えないかと常々思っている。MultiWOZの旅行対話では一応シミュレーターを作っているが、他のドメインでシミュレーターを作るとなると難しい。何か使える技術があったら使いたいと思っている。

社会実装・社会貢献や今後の展望

福島 最後にまとめとして、それぞれの領域で社会実装・社会貢献が進んでいるところ、これから狙っていくところがあると思うので、それに向けた課題や展望を一言ずつお願いしたい。今後CRDSで戦略提言を考えていく上で参考にさせていただきたいと思っている。

野田 私自身は交通系で実サービスも始めているが、意外といろいろなサービス連携の話が出てきている。そういう意味で社会に対するアウトプットもこれから増えていくと思う。サービス連携では、今日の話題にあった交渉・調整、多目的展開など、学術的にもいろいろ取り組める部分が広がっているという印象を持っている。

森永 私が取り組んでいる交渉を通して皆で調整していいところ見つけるとするのは、とてもネットワーク効果が効く建付けになっている。サプライチェーン全体の効率もドローンの空間使用率も、皆が使うようになるとご利益があるが、普及率の低い初期の間は、頑張ってもご利益があまりない。何らかの方法で、エイヤと普及率を上げるのがハードルであり、チャレンジになる。

山田 いろいろ方法論を開発して、わりときちんと評価して、論文発表しているが、なかなか企業で使ってくれない、うまくアウトリーチできないと感じていた。最近は、企業と共同研究を行ったり、技術顧問のような形で関わったりすることで、実際のサービスを作るときにアドバイスするというのが直接的に社会やビジネスにつながると感じている。書籍を出すというのかなり効果がある。研究者の世界では論文がとても価値があるが、それだけでは社会貢献は難しいと感じるようになった。

小野 非常に面白く聞かせていただいた。HAIはまだトイプロブレムの面があり、特に東中先生と駒谷先生がやられているような、実世界で使える対話エージェントに近づけたいと思う。また、最近は量が質を凌駕するという面も感じる。HAIは理想を追っている人が多いようにも思うので、もう本当に実際に使えるものをガリガリ作っていくことも考えてみたいと思った。

東中 いろいろ難しい問題に気付かされた。目の前の人間の発話をどう処理するかということばかりやっていたが、個人のモデル化とか、全体としてどう最適に振る舞うようにするかとかを踏まえて、対話システムを作っていくかといけないと思い直した。ムーンショットで複数の人間がインタラクションする中でどう対話を使っていくか、1人で10人・20人と対話するシステムを作るといったこともやっているが、対話システムのコミュニティーの中で、ベースとなる考え方がまだあまり持っていない。そのあたりから新しい取り組みが生まれるかもしれない。また、深層学習のコンピューティングパワーゲームに巻き込まれていて、日本ではLINEが頑張っていると思うが、どうすべきか悩ましい。

駒谷 今日は非常に勉強になった。HAIは自分の研究と深い関係があると思っていたが、これまであまり接点のなかったマルチエージェントシミュレーションも関係があると感じた。特に、学習するときにシミュレーションでデータを作るとは、重要なパラダイムになると思う。ユーザーモデリングで評価するというのも、融合によって新しい流れが出てくるのではないかと思った。また、機械学習は過去の再生で、過去にないものは全然できないし、パラメーターが増えるにつれて丸暗記に近づいていると思うので、全く未知のことをどう扱うか、本当に新しいことをどうやるのか、といった方向が今後のチャレンジに

なる。

伊藤 対話やマルチエージェントで新しい技術・応用の成功例が出てきたのは、深層学習によるところが大きい。そこでもう一つ重要なのは、人間と人間のコネクティビティー、ハイパワーコネクティビティー、瞬時につながっている状態がずっと続いているという状況だと思う。これは、民主主義や我々の社会システムができた頃にはなかった現象なので、そういったところを捉えた新しい社会システムを考えていかなければならないと思う。そのとき、マルチエージェントシステムやメカニズムデザインに加えて、エージェントと人間の関わりとして、HAI、対話エージェント、マルチモーダルインターフェースというものがとても重要になってくると思うので、一緒に何かできるのではないかなと思う。

荒井 人工知能研究において、マルチエージェントというキーワードは、伊藤先生のご意見に象徴されるように人間－人間、人－機械の相互作用のモデル化、支援（というより Serve）にフォーカスしすぎている気がする。社会システムの構成要素は、人、機械（人工物）だとすると、機械－機械の相互作用も重要になり、さらには、対自然、あるいは、自然との調和を包含することを考える必要がある。提携や交渉の先に実現される解が、どのレベルのシステムの目的関数（効用）を最大化しているのかといった視点が少し一部の人間中心に偏っている。SDGsを視野に入れるのであれば、時に人にとっては、必ずしも愉快的な解でなくとも、行動変容させ（説得）、社会的受容につなげる橋渡しができるような信頼性の高いシステム構築を目指す必要があると思う。さまざまな分野（AI、ヒューマンインターフェース、ロボティクス、認知科学、心理学、哲学など）の知見というのは、偏狭で、加えて自然科学、気候なども視野にいれた統合的シミュレーター構築へと広げないと真の社会シミュレーションとは言えないと思う。

大沢 大変有意義な機会だった。経済発展とか、SDGsとか、いろいろな考え方があるが、今これだけ複雑化した社会の中にあるさまざまな問題に対して、科学的な理論や技術的な発展をもって解決することしかできないと自分は思っている。人工知能の分野はその可能性を持っていて、今日の話題に限っても、その仕組みを作る取り組みや、シミュレーションによって将来を見通すことや、人間同士のインタラクションを本質的に深めることなど、人間社会の本質的な部分を相当含んでいて、そこに関わる問題を断片的にでも議論できたと思う。我々は人間社会に対して、もしくは、日本という国の発展に対して、どのように貢献できるのか、社会からどう見られるのかを常に意識しながら、自分の研究を相対化して進めていくことが重要だと思う。

付録

付録1 ワークショップ開催概要

開催日程：2022年1月26日（水）13：00～17：00（4時間）

開催形態：オンライン（Zoom ミーティングを使用）

アジェンダ：

- (1) 13：00-13：05 [5分] 開催挨拶
木村 康則（JST CRDS 上席フェロー）
- (2) 13：05-13：20 [15分] エージェント技術動向：俯瞰報告書2021年版から（発表）
福島 俊一（JST CRDS フェロー）
- (3) 13：20-14：00 [40分] マルチエージェント領域（1）
 - ・注目動向や今後の方向性（発表2件）
伊藤 孝行（京都大学大学院情報学研究科 教授）
荒井 幸代（千葉大学大学院工学研究院 教授）
 - ・討議
- (4) 14：00-14：40 [40分] マルチエージェント領域（2）
 - ・注目動向や今後の方向性（発表2件）
野田 五十樹（北海道大学大学院情報科学研究院 教授）
森永 聡（NEC データサイエンス研究所 上席主席研究員）
 - ・討議
- (5) 14：40-14：50 [10分] 休憩
- (6) 14：50-15：30 [40分] HAI（Human-Agent Interaction）領域
 - ・注目動向や今後の方向性（発表2件）
山田 誠二（国立情報学研究所 教授）
小野 哲雄（北海道大学大学院情報科学研究院 教授）
 - ・討議
- (7) 15：30-16：10 [40分] 対話エージェント領域
 - ・注目動向や今後の方向性（発表2件）
東中 竜一郎（名古屋大学大学院情報学研究科 教授）
駒谷 和範（大阪大学産業科学研究所 教授）
 - ・討議
- (8) 16：10-16：55 [45分] 総合討議
上記メンバーに加えて
大沢 英一（公立はこだて未来大学大学院システム情報科学研究科 教授）
- (9) 16：55-17：00 [5分] まとめ

聴講参加登録者（上記メンバー以外）：

文部科学省

4名 研究振興局 参事官（情報担当）付

1名 科学技術・学術政策局 研究開発戦略課 戦略研究推進室

経済産業省

1名 産業技術環境局 研究開発課

新エネルギー・産業技術総合開発機構（NEDO）

3名 ロボット・AI部

1名 技術戦略研究センター（TSC）

科学技術振興機構（JST）

10名 研究開発戦略センター（CRDS）

6名 戦略研究推進部

1名 未来創造研究開発推進部

1名 社会技術研究開発センター（RISTEX）

1名 経営企画部

計29名

付録2 参考文献・参考情報

第1章 エージェント技術：俯瞰報告書2021年版から

- [1] 科学技術振興機構研究開発戦略センター，「研究開発の俯瞰報告書：システム・情報科学技術分野（2021年）」，CRDS-FY2020-FR-02（2021年3月）。
- [2] 科学技術振興機構研究開発戦略センター，「人工知能研究の新潮流～日本の勝ち筋～」，CRDS-FY2021-RR-01（2021年6月）。

第2章 マルチエージェント領域（1）

- [1] D. C. Parkes and M. P. Wellman, “Economic Reasoning and Artificial Intelligence.” Science 349 (6245), pp. 267–272 (July 16, 2015). DOI: 10.1126/science.aaa8403
- [2] Stuart Russell and Peter Norvig, “Artificial Intelligence: A Modern Approach (4th ed.)”, Pearson Education Limited (2021). (18 Multiagent Decision Making)
- [3] V. Conitzer and T. Sandholm, “Complexity of mechanism design”, Proceedings of the 18th Conference on Uncertainty in Artificial Intelligence, pp. 103–110 (2002).
- [4] Paul Dütting, Zhe Feng, Hari Narasimhan, David C. Parkes and Sai S. Ravindranath, “Optimal Auctions through Deep Learning”, Proceedings of 36th International Conference on Machine Learning: ICML’19, pp. 1706–1715 (2019).
- [5] Roger B. Myerson, “Optimal auction design”, Mathematics of Operations Research, Vol. 6, No. 1, pp. 58–73 (1981). DOI: 10.1287/moor.6.1.58
- [6] Takayuki Ito and David Parkes, “Instantiating the Contingent Bids Model of Truthful Interdependent Value Auctions”, Proceedings of the Fifth International Joint Conference on Autonomous Agents and Multi-Agent Systems: AAMAS2006, pp. 1151–1158 (2006).
- [7] Dov Monderer and Moshe Tennenholtz, “Strong mediated equilibrium”, Artificial Intelligence, Vol. 173, Issue 1, pp. 180–195 (2009).
- [8] Noam Brown and Tuomas Sandholm, “Superhuman AI for heads-up no-limit poker: Libratus beats top professionals”, Science, Vol. 359, Issue 6374, pp. 418–424 (December 17, 2017).

DOI: 10.1126/science.aao1733

- [9] Noam Brown and Tuomas Sandholm, “Superhuman AI for multi-player poker”, Science, Vol. 365, Issue 6456, pp. 885–890 (July 11, 2019). DOI: 10.1126/science.aay2400
- [10] Todd W. Neller and Marc Lanctot, “An Introduction to Conterfactual Regret Minimization”, Tutorial (2013).
- [11] Martin Zinkevich, Michael Johanson, Michael Bowling and Carmelo Piccione, “Regret minimization in games with incomplete information”, Proceedings of the 20th International Conference on Neural Information Processing Systems: NIPS'07, pp. 1729–1736 (December 2007).
- [12] Marc Lanctot, Kevin Waugh, Martin Zinkevich, and Michael Bowling, “Monte Carlo sampling for regret minimization in extensive games”, Proceedings of the 22nd International Conference on Neural Information Processing Systems: NIPS'09, pp. 1078–1086 (December 2009).
- [13] Milind Tambe, “Security and Game Theory: Algorithms, Deployed Systems, Lessons Learned”, Cambridge University Press (2011).
- [14] Rubinstein, “Perfect equilibrium in a bargaining model”, Econometrica, Vol. 50, No.1, pp. 97–109 (January 1982).
- [15] Rubinstein, “A bargaining model with incomplete information about time preference”, Econometrica, Vol. 53, pp. 1151–1172 (January 1985).
- [16] Jeffrey S. Rosenschein and Gilad Zlotkin, “Rules of Encounter: Designing Conventions for Automated Negotiation among Computers”, The MIT Press (1994).
- [17] Allan Dafoe, Yoram Bachrach, Gillian Hadfield, Eric Horvitz, Kate Larson and Thore Graepel, “Cooperative AI: machines must learn to find common ground”, Nature, Vol. 593, No. 7857, pp. 33–36 (May 4, 2021). DOI: 10.1038/d41586-021-01170-0
- [18] Anita W. Woolley, C. F. Chabris, Alex Pentland, N. Hashmi and Thomas W. Malone, “Evidence for a collective intelligence factor in the performance of human groups”, Science, Vol. 330, Issue 6004, pp. 686–688 (2010).
- [19] マルチエージェント逆強化学習の分類と主要論文

	2-agent	N-agent
Single-stage		Waugh11
General-sum stochastic game	Lin18-1	Ziebart10 Common goal Natarajan10 Reddy12
Zero-sum game	Lin14-1 Lin14-2 Lin18-2	Sub-opt Wang18
Other		Swarm ŠoŠić17 Multi-Robot Bogert18
Ng-base	Bayes-base	MaxEnt-base
		Other

- [Lin14-1] X. Lin, P. A. Beling, and R. Cogill, “Comparison of Multi-agent and Single-agent Inverse Learning on a Simulated Soccer Example,” arXiv:1403.6822 [cs], Mar. 2014.
- [Lin14-2] X. Lin, P. A. Beling, and R. Cogill, “Multi-agent Inverse Reinforcement Learning for Zero-sum Games,” arXiv:1403.6508 [cs], Mar. 2014.
- [Lin18-1] X. Lin, S. C. Adams, and P. A. Beling, “Multi-agent Inverse Reinforcement Learning for General-sum Stochastic Games,” arXiv:1806.09795 [cs, stat], Jun. 2018.
- [Lin18-2] X. Lin, P. A. Beling, and R. Cogill, “Multiagent Inverse Reinforcement Learning for Two-Person Zero-Sum Games,” IEEE Transactions on Games, vol. 10, no. 1, pp. 56–68, Mar. 2018.
- [Wang18] X. Wang and D. Klabjan, “Competitive Multi-agent Inverse Reinforcement Learning with Sub-optimal Demonstrations,” in Proceedings of the 35th International Conference on Machine Learning, Stockholmsmässan, Stockholm Sweden, 2018, vol. 80, pp. 5143–5151.
- [Wangh11] K. Waugh, B. Ziebart, and D. Bagnell, “Computational Rationalization: The Inverse Equilibrium Problem,” in Proceedings of the 28th International Conference on Machine Learning (ICML-11), New York, NY, USA, 2011, pp. 1169–1176.
- [Ziebart10] B. D. Ziebart, J. A. Bagnell, and A. K. Dey, “Modeling Interaction via the Principle of Maximum Causal Entropy,” in ICML, 2010.
- [Natarajan10] S. Natarajan, G. Kunapuli, K. Judah, P. Tadepalli, K. Kersting, and J. Shavlik, “Multi-Agent Inverse Reinforcement Learning,” in 2010 Ninth International Conference on Machine Learning and Applications, Washington, DC, USA, 2010, pp. 395–400.
- [Reddy12] T. S. Reddy, V. Gopikrishna, G. Zaruba, and M. Huber, “Inverse reinforcement learning for decentralized non-cooperative multiagent systems,” in 2012 IEEE International Conference on Systems, Man, and Cybernetics (SMC), Seoul, Korea (South), 2012, pp. 1930–1935.
- [Šošić17] A. Šošić, W. R. KhudaBukhsh, A. M. Zoubir, and H. Koepl, “Inverse Reinforcement Learning in Swarm Systems,” in Proceedings of the 16th Conference on Autonomous Agents and MultiAgentSystems, Richland, SC, 2017, pp. 1413–1421.
- [Bogert18] K. Bogert and P. Doshi, “Multi-robot inverse reinforcement learning under occlusion with estimation of state transitions,” Artificial Intelligence, vol. 263, pp. 46–73, Oct. 2018.

第3章 マルチエージェント領域（2）

- [1] 「ウォルマートがコロナ禍に導入検討、発注条件ネゴシ契約するAI」, 日経クロストrend（2020年7月28日）.
- [2] 三宅陽一郎, 「大規模デジタルゲームにおける人工知能の一般的体系と実装—FINAL FANTASY XVの実例を基に—」, 人工知能学会論文誌, 35巻2号, pp. 1–16（2020年）.
- [3] 三宅陽一郎, 「デジタルゲームにおけるメタAI—キャラクターAI—スペーシャルAI 動的連携モデル」, 第34回人工知能学会全国大会（2020年）.
- [4] 三宅陽一郎・鳥海不二夫, 「MCS-AI 動的連携モデルがゲームデザインに与える影響の比較検証実験」, 第35回人工知能学会全国大会（2021年）.

第4章 HAI（Human-Agent Interaction）領域

- [1] HAIに関連する国際会議のリスト
 - ・ Int. Conf. on Human-Agent Interaction (HAI)

- ・ Int. Conf. on Human-Robot Interaction (HRI)
- ・ Annual CHI Conf. on Human Factors in Computing Systems
- ・ Int. Conf. on Intelligent User Interface (IUI)
- ・ Annual meeting of the Cognitive Science Society (CogSci)
- ・ Int. Conf. on Intelligent Virtual Agents (IVA)
- ・ Int. Conf. on Autonomous Agents and Multi-agent Systems (AAMAS)
- ・ Int. Symp. on Robot and Human Interactive Communication (RO-MAN)
- ・ Int. Conf. on Intelligent Robots and Systems (IROS)

[2] HAIに関連するジャーナルのリスト

- ・ Nature
- ・ Science
- ・ Proceedings of the National Academy of Sciences (PNAS)
- ・ Scientific Reports
- ・ PLOS ONE
- ・ Frontiers in Psychology/Robotics and AI – Human media interaction –
- ・ IEEE Access
- ・ ACM Trans. on Human-Robot Interaction (HRI)
- ・ Int. Journal of Social Robotics (SORO)
- ・ ACM Trans. on Interactive Intelligent Systems (TiIS)

- [3] Toyoaki Nishida, Atsushi Nakazawa, Yoshimasa Ohmoto and Yasser Mohammad, “Conversational Informatics: A Data-Intensive Approach with Emphasis on Nonverbal Communication”, Springer (2014).
- [4] Boris De Ruyter and Emile Aarts, “Ambient intelligence: visualizing the future”, Proceedings of the working conference on Advanced visual interfaces (2004).
- [5] Richard H. Thaler and Cass Sunstein, “Nudge: Improving Decisions About Health, Wealth, and Happiness” (updated edition), Penguin Books (2009).
- [6] Tetsuo Ono, Michita Imai and Ryohei Nakatsu, “Reading a Robot's Mind: A Model of Utterance Understanding based on the Theory of Mind Mechanism”, Proceedings of the Seventeenth National Conference on Artificial Intelligence: AAAI-00 (July-August 2000).
- [7] 小川浩平・小野哲雄, 「ITACO: メディア間を移動可能なエージェントによる遍在知の実現」, ヒューマンインタフェース学会論文誌, 8巻3号, pp. 373-380 (2006年8月).
- [8] オーガナイズドセッション「“ナッジ”エージェント: 人をウェルビーイングへと導くエージェントの構築へ向けて」, 人工知能学会全国大会 (2019年6月, 2020年6月).
- [9] Chenlin Hang, Tetsuo Ono, and Seiji Yamada, “Designing Nudge Agents that Promote Human Altruism”, Proceedings of the 13th International Conference on Social Robotics: ICSR 2021 (November 2021). arXiv: 2110.00319
- [10] Ana Caraban, Evangelos Karapanos, Daniel Gonçalves and Pedro Campos, “23 Ways to Nudge: A Review of Technology-Mediated Nudging in Human-Computer Interaction”, Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems: CHI’19 (May 2019). DOI: 10.1145/3290605.3300733
- [11] 科学技術振興機構研究開発戦略センター, 「俯瞰セミナー&ワークショップ報告書: トラスト研究の潮流 ~人文・社会科学から人工知能、医療まで~」, CRDS-FY2021-WR-05 (2022年2月).

第5章 対話エージェント領域

- [1] 東中竜一郎・他, 「対話システムライブコンペティション4」, 第12回対話システムシンポジウム論文集 (2021年11月).
- [2] 東中竜一郎, 「私のブックマーク: 対話システム」, 人工知能学会誌, 36巻5号, pp. 644-650 (2021年9月).
- [3] 東中竜一郎・光田航, 「対話システムでは今何が問題になっているのか?」, 情報処理学会誌, 62巻10号, pp. e24-e31 (2021年10月).
- [4] Seungwhan Moon, et. al., “Situated and Interactive Multimodal Conversations”, arXiv: 2006.01460 (2020).

URL 情報について

本文・脚注に記載の URL は 2022 年 3 月時点のものである

俯瞰ワークショップ報告書

CRDS-FY2021-WR-11

エージェント技術

令和 4 年 3 月 March 2022

ISBN 978-4-88890-783-5

国立研究開発法人科学技術振興機構 研究開発戦略センター

Center for Research and Development Strategy, Japan Science and Technology Agency

〒102-0076 東京都千代田区五番町 7 K's 五番町

電話 03-5214-7481

E-mail crds@jst.go.jp

<https://www.jst.go.jp/crds/>

本書は著作権法等によって著作権が保護された著作物です。

著作権法で認められた場合を除き、本書の全部又は一部を許可無く複写・複製することを禁じます。

引用を行う際は、必ず出典を記述願います。

This publication is protected by copyright law and international treaties.

No part of this publication may be copied or reproduced in any form or by any means without permission of JST, except to the extent permitted by applicable law.

Any quotations must be appropriately acknowledged.

If you wish to copy, reproduce, display or otherwise use this publication, please contact crds@jst.go.jp.

FOR THE FUTURE OF
SCIENCE AND
SOCIETY



<https://www.jst.go.jp/crds/>