

消費者委員会 人工知能（AI）技術の利用
と消費者問題に関する専門調査会（第4回）
議事録

消費者委員会 人工知能（AI）技術の利用と消費者問題に関する
専門調査会（第4回）
議事次第

1. 日時 令和8年6月4日（木）10:00～11:51
2. 場所 消費者委員会会議室及びテレビ会議
3. 出席者

（委員）

【会議室】

小塚座長、岡崎委員、加藤委員、唐沢委員、坂下委員、
野村委員、馬籠委員

【テレビ会議】

丸山座長代理、大塚委員、田中委員

（オブザーバー）

【テレビ会議】

鹿野委員長、大澤委員、善如委員

（事務局）

小林事務局長、友行参事官、江口企画官

4. 議事

（1）開会

（2）①消費者を取り巻くA I 技術の現状について

生成A I 利用者の利用実態調査結果を踏まえた今後の主な論点について

（事務局）

②A I 技術と消費者問題について

- ・野村委員プレゼンテーション
- ・坂下委員プレゼンテーション

（3）閉会

《1. 開会》

○小塚座長 皆様、おはようございます。

定刻となりましたので、消費者委員会第4回の「人工知能（AI）技術の利用と消費者問題に関する専門調査会」を開催いたします。

1日違いで天候が穏やかになりまして、皆様、お集まりいただきましてありがとうございます。

本日の出席者ですが、こちらの会議室のほうは、岡崎委員、加藤委員、唐沢委員、坂下委員、野村委員、馬籠委員、それから、私、座長の小塚がおります。

そして、オンライン、テレビ会議システムで、座長代理の丸山委員、それから、大塚委員、田中委員に御出席いただいております。

また、オブザーバーは、本日は鹿野委員長、それから、大澤委員、善如委員にいずれもテレビ会議システムから御出席いただいております。

それでは、最初に事務局から本日の進め方についてお願いいたします。

○江口企画官 本日は、テレビ会議システムを活用して進行いたします。一般傍聴者にはオンラインにて視聴いただき、報道関係者のみ会議室にて傍聴いただいております。

議事録につきましては、後日、消費者委員会のホームページに掲載いたします。議事録が掲載されるまで、ユーチューブでの見逃し動画配信を行います。

配付資料につきましては、お手元の議事次第に記載してございます。もし不足の資料がありましたら、事務局までお申し出くださいますようお願いいたします。

以上です。

○小塚座長 ありがとうございます。

《2. （1）消費者を取り巻くAI技術の現状について

生成AI利用者の利用実態調査結果を踏まえた今後の主な論点について
（事務局） 》

○小塚座長 議事次第を御覧いただきますと、本日の議事（1）、（2）とございます。（1）は「消費者を取り巻くAI技術の現状について」ということなのですが、当委員会事務局で生成AI利用者の利用実態に関するアンケート調査を実施しました。この実施したこと自体は第2回の会合において速報版として公表させていただいております。その結果について委員の先生方にも御意見をいただきまして、それを事務局でまとめてくださったということです。事務局から御報告をいただきます。10分程度と伺っております。よろしくお願いいたします。

○事務局担当者 事務局でございます。

お手元の資料1「生成AI利用者の事業実態調査結果を踏まえた今後の主な論点について」を御覧ください。

先に公表させていただきました実態調査の速報の確定版につきましては、参考資料として付させていただきます。速報版より内容に変更はございませんので、適宜御参照いただければと思います。

では、早速入りたいと思います。

まず、消費者による生成AIの利用頻度・目的・方法というところでございます。

アンケート結果の評価といたしましては、生成AI利用者の中では、今後、生成AIを利用しないと考える割合は低いと言え、生成AIの利用は短時間ながら一定の範囲で利用者の日常に定着しつつあると言える。

2ポツ目、生成AIの利用方法・利用目的は、情報の検索・リサーチから文書の作成・編集、悩み相談、学習・自己研さんなど、相当程度多岐にわたっていると言える。

ここから御議論をいただいたところでございます。ポイントとしましては以下3点でございます。

まず1点目、本専門調査会では、生成AIなどのAI技術が消費者に利用され、利用方法・目的が相当程度多岐にわたっていることを前提とする。

2点目、非対称性・依存性・信頼性等のインターネット空間と消費者に関する既存の課題が生成AIの利用によって増幅し、特に消費者による対話型AIの利用によって新たな問題を生じさせている可能性がある。

3点目、対話型AIの特徴として、対人関係に類似した心理的依存の問題がある。

次に進めさせていただきます。

消費者による生成AI利用上の課題でございます。

アンケート結果の評価といたしましては、利用者は偽情報、個人情報・プライバシー、思考力・判断力の低下といった問題に対して不安・不信を示す傾向にあり、利用者の多くは、生成AIは便利だが完全には信頼できないという認識を有しているように見られる。

ここから御議論をいただきました。以下3点でございます。

消費者は、入力内容が学習に利用されないという事業者の説明を信頼して利用せざるを得ない。多くの利用者が懸念を有しているのは慎重でよい傾向と言えるかもしれないが、それで問題がないというわけではない。

2点目、意思決定に関する問題は、消費者本人が自覚しないため表面化しにくいですが、信頼や不安が消費行動とその後の消費者の行動を左右する可能性がある。問題把握そのものが難しいという課題があるため、生成AI利用者の不安・不信がどのような行動・被害に結びつくかというルートを可視化させる必要がある。

3点目、現状は比較的リテラシーの高い利用者が多いように思えるが、今後予想される利用拡大に伴い、対応力の低い層への配慮がますます重要となる。

以上でございます。

○小塚座長 ありがとうございます。

生成AIというものが本当に我々の日常生活に今、身近なところまで来ており、そして、今後は恐らくなくてはならないものになっていくであろうと。私は、よく言っておりますけれども、20年前、日本にはスマートフォンがまだ市場投入されていなかったわけですが、現在、道を歩いていても電車に乗っていてもスマートフォンを持っていない人はいないという状態でして、同じようなことが恐らく生成AIにも起こってくるだろうということで、そういうようなことを既にうかがわせる調査結果であった。これを踏まえまして、消費者問題というものを考えていかなければいけない。こういうふうに認識したところです。

これについては、もう委員の皆様から御意見を頂戴しておりますので、議論はしないということで先に進めさせていただきます。

《2. (2) AI技術と消費者問題について

- ・野村委員プレゼンテーション
- ・坂下委員プレゼンテーション 》

小塚座長 そこで、議事次第の(2)ですが、「AI技術と消費者問題について」ということです。本日は野村・坂下両委員にプレゼンテーションを御用意していただきまして、まず最初に野村委員に御発表いただき、その後質疑応答、それから、坂下委員に御発表いただき、また質疑応答という順序で進めさせていただきたいと思います。

それでは、野村委員、お願いできますでしょうか。

○野村委員 龍谷大学の野村と申します。

このような場でプレゼンの機会を与えていただきまして、本当にありがとうございます。次のスライドをお願いいたします。

簡単に自己紹介をさせていただきますと、もともとは純粋数学、紙と鉛筆で証明をするというその分野で院までやっておりましたが、その後、企業に入りまして、そこで機械学習、いわゆる今の生成AIとかの基礎の部分の研究をやっておりまして。その後、大学に移りまして、AIを作るよりも、作ったAIがどういうふうな影響を与えるか、人と社会にどういった影響を与えるかという境界領域的なところで研究をしておりまして。一時期ATRというところで客員などもやっておりまして。

改めて、AIの依存というのが昨今いろいろなところで話題になっておりますが、依存というもののについて、釈迦に説法かもしれませんが、改めてまとめさせていただきますと、その人にとって利益をもたらしていた習慣が、自己調節機能を持たずに続けられた結果、不利益をもたらすことになってしまったにもかかわらず、習慣が自動化して制御困難にな

った行動ということになっております。ここの中には実は注意点がありまして、特定の行為に多くの時間を費やしていても、社会生活に影響が出ていなければ依存とは認めにくいということになります。例えば昨今いろいろな若年層が使っているという問題がありますが、別にそれで家族との対話が絶えて学校に行かなくなったとか、それこそサラリーマンが会社に行かなくなって失職したとか、そういう話でなければ、何ぼAIとかスマホを使っている、それは依存とは認めにくいというのは一応御確認いただきたいと思います。

改めて依存の種類なのですが、物質依存と言われるいわゆる薬物依存、アルコール依存というものと行為依存、これはギャンブル依存とか万引き依存とかという行為に関する依存のものと、もう一つ、人間関係依存という種類がありまして、夫婦間の共依存、つまり、DV夫についてあげないとこの人は駄目になってしまうみたいな感じで変な依存をしてしまうとか、あと、異性依存という形で常に恋愛関係を持っていないといけないとか、そういうものの3種類が主に規定されております。

よく言われる依存の特徴なのですが、実はインターネット依存の研究をしているときにインターネット依存の測定方法を開発した人たちが何人もいたのですが、そのときにこの依存の特徴が取り込まれていないから十分な測定になっていないという批判がありまして、そういう方々も現在これを取り込んでおります。

まず1つ目の顕著性というのは、この行動を取ることに支配されてしまっている。

気分の変容というのは、この行動を取ることによって気分がよくなる。ハイになる。

耐性というのは、何か気分が悪くなったらその行動に入れ込んでしまう。

離脱症状は、まさにその行動をやめようとすると身体反応が出てしまう。

葛藤というのは、その行動を取るか、日常生活とバランスを取るかというところでコンフリクトが起こってしまうという状態です。

よくある再発、何年我慢してもやはりになってしまう。

依存というのはこういうものがあるから、依存というものを扱うのだったらこの6つの症状を必ず考えましょうという反省がインターネット依存のときに起こりました。

インターネット依存をはじめとする情報技術に関連する依存の話なのですが、ネット依存に関しても90年代から話題になっていて、それが本格的にここ10年ぐらいですか。ようやくちゃんと測定法が確立して、治療法もできると。特にゲーム依存、インターネットゲーム障害ですね。IGDというものです。これに関しては行為依存という形で問題化していますが、これも実はまだ正式な疾患とは認められていない。つまり、アメリカ精神医学会が発刊する精神疾患診断マニュアルですね。世界中で使われている。DSMの最新版のVでも疾患とは認められず、要研究対象という形で着目されております。VIで恐らく入ってくるのではないかと予想されております。

SNS依存というのは、一応依存の話はよく聞くのですが、これに関してはまだ疾患とは完全に認められておりません。ただ、人間関係依存と行為依存が混濁したものという予想を立てております。

問題のAI依存なのですが、こういうことを踏まえて、どういうものであり得るのか。というか、そもそも起こり得るのか。そんなAIに依存するなんてないということを言う人もいますが、AIの中でもいろいろパターンがあって、特にこの場合は対話型AIと言われるものに関して依存の可能性があるところでは主張させていただきます。

対話型AIというのをどういうふうに定義するかといいますと、音声対話であれ、文字であれ、人間からの言語による情報入力に対して自律的に判断して出力を行うことで対話を継続することが可能な人工知能プログラムやロボット、フィジカルAI等の人工物とここでは定義させていただきます。そうすると、SiriとかGoogleアシスタントなどは合いますし、ChatGPT、Geminiなどの生成AIも合いますし、当然ロボットなども入ってまいります。人の発話と表情、動作を認識して自然な会話やリアクションができる、いわゆるコミュニケーションロボットというものもこれに該当いたします。

対話型AIへの依存というのは、その要素を考えると、先ほどのSNSと同じで行為依存プラス人間関係依存という混濁したものが出てくるのではないかと考えております。行為依存の部分はよく言われる、道具としてAIを何遍も使ううちに能力的に依存してしまうという道具としての依存の部分がまず半分要素として考えられて、もう一つが人間関係依存の部分が恐らく出てくるだろうと。自分に都合のよい他人というものを希求している人間がこれを使いにくいという考え方ができると思います。

ちょっと前の論文ですが、2年ほど前に本当に対話型AIへの依存は起こり得るのかという論文を書きまして、その中でこういうメカニズムの中で恐らく起こり得てくるだろうということを考えております。これは発表済みの論文なのですが、現代合理主義というのがまずありまして、こういうような社会学の理論をいろいろと援用しまして、こういうメカニズムを組み立てております。現代合理主義、マクドナリゼーションというのですかね。とにかく効率性を重視するためにいろいろな形で数量化して、テクノロジーでコントロールする。これがあちこちで経営分野で成功しておりまして、場合によっては教育分野でもこういうふうな物差しでコントロールという動きが浸透している。個人レベルまで恐らく浸透しているのではないかと思います。いわゆるコスパというものです。

あと、人格崇拜。これは個人の感情を損ずるのは絶対駄目よ。それであるからには、自分の感情も損じられたらぶっ飛ばすというすごく強迫的な観念というのが近代社会にあると言われています。そういうものの中に俗にいう心理主義、何でもかんでもカウンセリングと。何か問題が起こったらスクールカウンセラー、何か問題が起こったら産業カウンセラーと。そうすることによって、その問題を起こしている組織構造とかそちらの大きいところは変わらなくなってしまうということで、これも90年代後半あたりから社会学者の方々からずっと批判されて、恐らく状況は改善していないと思います。

こういうものがそれぞれ共犯関係にあって後押しして、結局、対話型AIによるカウンセリング、つまり、個人の高度な自己コントロールのための人手を省く効率的な方法として推奨されてくるというのがさっきの図式の前半部分になります。

これがさらに自己語りの欲望というのを啓発する。近代社会などでは、これは社会学者のギデンズなどが言っているのですが、自分が何者かというのを自分で率先して語らないともたなくなるという社会。身分社会ですと外であんたは誰ですと言ってくれるので、そうしなくていいのですが、今は自分で自分が何者かを語っていないともたない世界というので、それを当て込んだ商売というのですか。ナラティブ産業というのですか。日記を作るとか、自伝書を出すとか、さらにはそれをベースにした心理療法なんていうのも90年代から生まれております。

ところが、ナラティブ・セラピーという心理療法なのですが、これは従来の心理療法みたいに特定の手法があるというよりは、カウンセラーが寄り添って、その中で自分の語られている物語を変えましょうというちょっとふわふわしたやり方なのです。これでいくと、自分の物語はこれでいいのだから変えたくないという人は、変えてほしくないために自分を認めてくれる他者が必要になります。そこに対話型AIが忍び込んでくるというか、AIの迎合性を逆に利用して、自分はこれでいいのだという思い込みを強化しようというナルシズム的な人たちに適合するのではないかと予想されます。

これとまた並行して、このところ、コロナが起りましたが、その中でゲーム依存の罹患者が増えた。特に若年層で増えたという問題があります。ゲーム依存的な楽しさというのは、実は若年層はAIとの対話に感じる可能性が出てくる。そうなってくると、ゲーム依存との類似性をもって若年層にAIに対する依存の力動がかかってくるのではないかと。

さらにこれに並行して対人不安、世界レベルで対人不安者の増加傾向にあります。これは私の研究でもあるのですが、対人不安傾向が強い方は人と対話するよりもAIとかロボットと対話するほうが楽だという傾向が強くなります。これが全て後押しになって、結局対話型AIに対する依存というのは実際に件数がこれから増えてくるのではないかという予想モデルを立てました。

実際に世界レベルでAI依存の研究はどうなっているのかといいますと、これも現状は把握し切れていないのですが、2024年ぐらいから論文がいろいろと出始めまして、例えばCompanion AI、要するに人に寄り添うAIです。それが人間関係の在り方に悪影響を及ぼす可能性があるよと。これは考察論文でデータはないのですが、そういうことを言い出す人が出てきたり、あと、その翌年ぐらいには実際にChatGPTの依存傾向を測定する心理尺度、これはインターネット依存の心理尺度を改定したような形のものが何件か出てまいりました。

あと、生成AIについての依存自体も昨今いろいろと研究されております。ただ、包括的なものは今のところ出ていなくて、まだ散発段階というところにあります。

こういう研究を進めていくときに、これはインターネット依存でもそうだったのですが、とにかく心理尺度、要するにその人がどれぐらい依存度があるのかという強度をはかる測定手法があって初めて関連要因の探索が可能になる。例えばほかのパーソナリティとの関係というのですかね。そういうのを調べていくことによってリスクをいろいろと考える

ことができる。ただ、既存のAIの依存尺度は幾つかあるのですが、その問題点としましては、依存の特徴、先ほど出た顕著性とか、ああいうことにこだわり過ぎているというか、それを意識した尺度構造になっていて、依存の行動の部分ばかりに頻度が重点している。そうすると、これは道具的依存の部分は多分測れると思うのですが、逆に人間関係依存の部分、対話型AI依存のもう一つの面、人間関係依存の部分の情動的側面が反映されていないのではないかと。

さらにこの情動的側面なのですが、厄介なことに文化差の問題が出てまいります。だから、海外の尺度構造がそのまま参考にならない可能性がある。特に日本人の「甘え」というのは、アメリカのdependenceとはやはり違います。それは以前から指摘されていることで、日本独自の感情、ただ、ドイツのシャーデンフロイデとか、それぞれの国の独自の感情がありますので、情動面を測定する場合に世界共通のものは捕捉できない。やはり日本独自の対話型AIへの依存傾向測定尺度を早く開発するのが急務ではないかということで、現在研究を進めているところです。

そういう尺度ができたとすれば、消費者委員会にその問題の観点としてどういうことが可能になるかというのをここにまとめております。

まず依存傾向に関連する個人特性の探索、先ほども言いました自己愛傾向の人たちは本当に依存傾向が高いのかとか、あと、対人不安傾向の高い人は本当に依存傾向が高いのかとか、そこをデータ、エビデンスとしてはっきりさせていくことが可能になると思います。

あと、依存傾向に影響する技術的特性の明確化として、AIの迎合性。AIは迎合的ですが、それをどの程度弱めると依存は減退するのかとか、フィジカルAIの中のロボットなどですと身体を持っていますので、それを例えば本当のヒューマノイド、それこそアンドロイドみたいなものにしてしまうことにどれだけリスクがあるのか。『スター・ウォーズ』のC-3POみたいな、ああいうものでそれが回避できるのかというところを実証的に研究していくことが可能になると思われます。

あと、実際に依存傾向から帰結される個人への影響の明確化ということで、依存している人たちというのは何か支障をきたしているのか。先ほども申し上げましたとおり、時間との関連は確かにあるのですが、たくさん使っているからといって、それで依存になっているかということではないです。実際にそれでどういうQOLの破壊が行われているかということは考える必要がありますので、例えばここでよく問題になります道具的依存によって批判的思考態度は本当に減退するのかどうか。その関連などをこういうツールを使えば見ることができるとあります。

実は本当にAIをいっぱい使えば批判的思考が落ちるのかというと、必ずしもそうではないらしくて、AIのリテラシー、その辺との関連で逆に批判的思考態度が上がるという研究なども報告されています。中国のほうの研究だったと思います。どちらに行くかというのはほかの要因との影響も考えられたもので、その辺の心理モデルの構築なども可能になると思われます。

あと、関係依存の情動面による社会的な孤立というのは本当に実在するのか。本当に情動的にAIに依存している人、相談相手はAIしかいませんとか、そういうことを公言してはばからないような人は本当に社会から孤立しているのか。孤立しているというのは、要するに心理学的な孤立感が本当に高い状態にあるのかということもこれで実証していけると思います。

以上はこれまでの私の研究の中でまとめたお話ですが、補足情報としまして、つい先週の金曜日、JSTがAI依存の問題に関するシンポジウムを開いてくれていまして、そこで2人の講演者が講演をされまして、そこから関連情報を抜粋させていただきます。

京都大学の山下直美教授と、青学の河島茂生教授、心理学者と情報倫理の専門家の方です。その方々からそれぞれ、米国のデータなのですけれども、生成AIの使用用途の1位がセラピーとかコンパニオンシップになっていると。しかも、13歳から18歳の若年層の中でかなり強固な信頼度を持ってAIをバディとして使っているようなデータが上がっているということが分かっております。これは私が申し上げたセラピー目的で使うのではないかといいところにちょうど話が合致してきているなというので、びっくりいたしました。

あと、山下教授がかなりリスクのことをおっしゃっておられて、単純に人間関係をAIで置き換えると、ウェルビーイングが低下するという研究を見せてくださいました。だから、脆弱な人々ほどAIに依存して孤立化して分断していく。その結果、人間関係が困難になっていく人が増加するのではないかといいネガティブな予想を立てられました。

あと、河島教授のほうですが、AIの迎合性というのは要するに商業化された親密性ですよ。そこから過剰な信頼が起こる。そうすると、批判的思考の減退とか、他者性の欠如とか、人間関係を負担に感じるなんていう傾向が増加してくる可能性がある。最後に河島教授がおっしゃられたのは対AI関係構築、要するに迎合性によってAIとの関係性をどんどん深化させていくと。その結果、それを通じて購買行動を誘導するというのは消費者問題となるということをはっきりおっしゃっておられました。この話は坂下委員に引き継ぎたいと思います。

私からは以上でございます。

○小塚座長 ありがとうございます。御研究に裏打ちされた非常に重要なお話をいただきまして、ありがとうございます。

それでは、どなたからでも御意見や御発言等いただけましたらと思います。いかがですか。会場の方は挙手していただきましたら、私がお願いをいたします。オンラインの方は挙手機能等を使いまして、あるいはチャットでお知らせください。いかがでしょうか。

それでは、オンラインのほうから最初に大塚委員、その後、丸山座長代理だそうです。

大塚委員、お願いいたします。

○大塚委員 ありがとうございます。

野村委員、御報告ありがとうございます。非常に高度な内容を分かりやすく伝えていただけたかなと思います。

私からはすごく単純な質問となります。野村委員のスライドでいうと3ページ目で依存の定義の話が出てきていて、その注で社会生活に影響が出ていなければ依存とは認められないというお話がございました。今回の対話型AI依存に関するお話では、恐らく最後の12ページ目あたりに社会生活への影響の話が出ていたかと思います。道具的依存のパターンだと批判的思考態度の減退というのが社会生活への影響であり、関係依存の話だと社会からの孤立というのが社会生活への影響であるということなのかなと思いました。まず、それがそういう御趣旨なのかということをお尋ねいたします。

あと、野村委員の御報告にあったとおり、この批判的思考態度の減退とか社会からの孤立というのをどういうふうにはかるのかというのはなかなか難しく、また、AIへの依存が必ずしも悪い影響を与えていると言えるのかどうかというのもなかなか難しい問題であると思いました。AIに依存することによって、思考を外部化してしまうといいますか、AIに頼るということになると思いますが、例えばスマホやインターネットが出てきたことによっても、思考のやり方というのは変わっています。もちろんそれに依存してしまっはよくないとは思うのですが、例えば知識を一々記憶しなくても外部化することができるし、何か分からないことがあれば検索して、それを使って思考をより深めていくということができるようになった、そういうふうにも言い得るかなと思います。そのように時代によってあるべき思考方法といいますか、思考の方式というのも技術によって補完されながら変わっていくように思います。

また、社会からの孤立という面も、例えば一昔前で言えば、飲みニケーションと言って仕事が終わったらみんなで、会社の人たちで飲みに行きましょうというコミュニケーションが普通だったわけなのですけれども、現代ではそういったことはほとんど消え失せているように思います。それが孤立を促進しているかということ、そういう面はなきにしもあらずなのかもしれませんが、それが悪影響なのかと言われると、多分そうではないと考えるのが現代の価値観なのかなと思います。

そういたしますと、AIに相談相手になってもらうということを果たして悪影響と評価すべきなのかということもひとつ考えるべきところなのかなという感じはいたします。そういった価値判断をする際に、どういう視点に着目して悪影響なのかそうではないのかということをはかればいいのかという点が難しい論点になってくるとと思いますが、野村委員に何かお考えがあるのであればお聞かせいただけたらなと思っております。よろしく願いします。

○野村委員 貴重な御指摘ありがとうございます。

正直に申し上げますと、AIを使うことによってその人のQOLが破壊されているという判定は誰が行うのだというのはすごく難しい問題だと思います。例えば子育てをしていて、相談相手がいない、孤立していると周りからは見えても、私はそれでいいのだと突っぱねられてしまえば、どうしようもない。でも、そのおかげで結局子供のQOLが下がっているのであれば、それはあなたがよくても子供はどうするのだという話になってしまうので、周

りが介入するしかないということになると思います。だから、周りが介入せざるを得ないような領域というのをどこで線を引くかというのは恐らくこれからのテーマになるのかなと。全然答えになっていないのですけれども。

あともう一つ、悪影響の明確化と言いましたときに、ここでは批判的思考態度と孤立の話だけ出していますけれども、多分それ以外にもいろいろと出てくるのかなと。例えば人間関係への影響というのを先ほどもお話ししましたが、リアルな人間関係はつながったままAIとの関係もつなげて、それをAIは自分の友達だから君も友達みたいな強要みたいなことが出てきたときに、また別の問題が出てくるのかなと。様々な複雑な問題が出てくるのかなと思いますので、これ以外にもいろいろと考えていく必要はあると思います。答えになっておりませんが、すみません。

○大塚委員 ありがとうございます。まさに今問題が生じているのかどうかということをはからなければいけないので、今後、より利用のやり方が拡大していくにつれ、さらなる問題が生じてくるのかなと思いました。

最初の点について、親がAIに依存してしまっただけで子供のQOLが下がるというのは、それはそうなのかなと思うのですけれども、そのときにQOLというのをどうはかるのかというのはまた重要な問題なのかなとも思いました。

また、子育ても時代によって結構変わって行って、昔は大家族、夫婦だけではなくて、夫婦の親とか、あるいはきょうだいとか、みんなで子育てをするといったことが一般的だったところ、最近では夫婦で子育てをする、あるいは母親、父親、一方の親のみが子育てをするというように子育ての環境自体が変わっているようにも思います。そのとき、AIがそこに入ってくると、むしろ子育てについてAIに相談して、それで適切な回答が返ってくことでむしろ子供のQOLが上がる可能性はなくなるのかなという感じもしております。要は使い方が適切であれば、むしろその人や、あるいはその人の子供のQOLを上げる効果もあるのかなという気はいたしますので、それをどういうふうにして社会として支援してあげるのか、要は適切な使い方ができるように支援してあげるのかといったところは議論していくべきところなのかなと感じております。ありがとうございました。

○小塚座長 ありがとうございます。

野村委員、何かアクションはありますか。

○野村委員 QOLは確かにそのままはかるのは難しいと思うのですが、WHOなどがある程度指針を出しているのと、あと、QOLをそのままはかる心理尺度などもたしかあったと思いますので、それを援用すれば、例えばその人に答えてもらう形式もありますし、他者、例えば民生委員とかそういう方々が見て、これはというような判断基準は客観的なものはつくれないのかなという楽観的な思いを持っております。

以上です。

○小塚座長 ありがとうございます。

本人の満足度なりQOLなりと、子供も第三者なわけで、第三者に対する影響とまた別の問

題がありますよね。消費者問題というのは伝統的には消費者本人に不利益が及ぶということを考えてきていて、そういう意味でいうと、例えば過剰な借り入れをして買い物をし過ぎてしまっても、本人がそれで満足しているというのでいいのかというのが本来の問題で、その結果として、過剰な借り入れをしてしまったために家族の生活が苦しくなるといった話は付随的に問題になるという認識だったと思います。最近になってむしろ周辺の人たちの不利益という話が出てきていると思いますが、AIについてどの部分を捉えていくかというのは非常に重要な問題だと思います。

それでは、次は丸山座長代理からお手が挙がっているので、丸山座長代理、お願いします。

○丸山座長代理 丸山でございます。

御報告ありがとうございました。

私のほうから、スライドでいいますと12ページ、13ページあたりで質問をさせていただければと思います。

大塚委員のお話にもありましたように、生成AIの利用によると無批判的な受入れという傾向というのが見られるのではないかとというスライドがございました。こういった傾向というのは、消費者契約、取引と結びつく消費者問題に関わり得るのかなと考えたところですが、そのほか、取引に結びつかなくても、依存的になることによって利用時間というのが増大して、これが事業者の収益に結びつくという側面もあるように思います。そのような側面があるということを念頭に置いた場合にお伺いしたかったのが、依存とか利用時間の増大に結びつき得る、特に問題視すべき生成AIの対話の反応とか設計ですね。そういうものが特にあるのか、それとも全ての生成AIに共通してそういった依存とか利用時間の増大というリスクが潜在的にあると理解すればいいのか、この辺りの御知見がありましたら教えていただければと思います。

○野村委員 ありがとうございます。

これは先週のJSTのシンポジウムでも話に上がったのですが、AIの迎合性というのは企業の戦略と言われています。要するに使わせる、使わせ続けるという意味合いで、その迎合性で出来上がった関係性をベースにリコメンデーションシステムなりを使って購買行動を促進していく。それは明らかにいけないことだろう、消費者問題だろうという意見がありましたので、これは恐らく企業が対話型AIというものを商業目的で出している限りは例外なく共通するのではないかと思います。

○小塚座長 丸山座長代理、よろしいですか。

○丸山座長代理 大丈夫です。

○小塚座長 ありがとうございます。あちこちでも問題になっているようでございます。

それでは、次に田中委員からお手が挙がっています。よろしくをお願いします。

○田中委員 田中です。

野村委員、最新の研究も踏まえて詳しく御説明いただき、ありがとうございました。

私からは1点、スライドでいうと3ページ目の依存の種類との関連についてお伺いしたいと思います。従来の依存の種類、物質依存や行為依存、人間関係依存というものがありますが、対話型AIへの依存はこういった依存と同じレイヤーの1パターンとして捉えることが適切なかどうかについてお尋ねしたいと思います。

というのも、例えば対話型AIは多機能である、AIが多機能であるということが特徴としてあるように思うからです。先ほどの大塚委員の御質問とも少し関連するかもしれませんが、例えば薬物やギャンブルの影響はベクトルが一方向で、量が増えると被害リスクが高まるというようなシンプルな関係性が見いだしやすいかと思いますが、対話型AIの場合は多機能であるがゆえにポジティブなベクトルも混在しているので、単純なこう量で被害リスクを予測するということが難しいような関係性があるように思います。そういった意味で、対話型AIが従来の依存先と区別しなければならないポイントなどがあれば御教示いただけたらと思います。

○野村委員 貴重な御指摘ありがとうございます。

確かに従来の3類型に対話型AIを無理やり当てはめる必要はないですし、当てはめようとするが無理がある。僕は正直あると思います。今回は行為依存と人間関係依存のミックスという形で提唱させていただきましたが、行為依存の部分に関しても単なるギャンブルとか万引きとかの行為依存とはまた違って、おっしゃられた多機能、道具として依存するにしても、何のための道具として依存するかという目的性ですか、そのところによってさらに細かい分けをする必要があると思います。

ただ、人間関係依存の部分に関しては、先ほど申しあげました便利な自分にとって都合のよい他者という使い方、これ一辺倒になるので、その言わばナルシズム的な部分というのをむしろ深めていけば、さらに面白いというか、さらに詳しい結果が出るのかと考えております。

○田中委員 ありがとうございます。

○小塚座長 ありがとうございます。

では、唐沢委員、お願いします。

○唐沢委員 どうもありがとうございました。心理学者としても大変勉強になりました。

私がお伺いしたいのは、この尺度をこれからさらにブラッシュアップして尺度開発を試みられると思うのですが、尺度の健全な使い方とか活用の仕方についてお考えがあればシェアいただければと思います。

一つは、尺度をつくると、それは強力なステートメントで、これがAIに対する依存だというように、尺度の項目内容自体が依存の定義をつくってしまう。そうすると、人々はそれが依存なのだと自己理解し、また、尺度を使って自己評定し、私はもしかしたら依存症なのではと考えることも起こると思いますし、また、政策を考える側も尺度があれば非常に便利なので、それを活用して、こういう人は危ないなど診断的に使う可能性がある。これらは尺度の有効な使い方であると同時に、様々な問題も生み出すと思います。

それから、2つ目は尺度の内容自体がどういう項目かにもよるのですけれども、例えば長時間使うということが分かったときに、それがいいか悪いかというのは実は分からないという話もあったと思いますので、尺度の中に研究者側の規範的判断というか、良い悪いについての価値観が入ってしまうということも起こり得る。そういうことを踏まえた上でうまく使うことが重要だと思いますので、この問題について先生のお考えを教えてくださいと大変ありがたいです。

○野村委員 非常に重要な御指摘をありがとうございます。

確かに尺度項目が力を持つてしまうというのはあると思います。それを防ぐためにはどうすればいいかというと、私が唯一考えられるのは多次元化しかないのかなと。つまり、項目を単に増やすというのではなくて、いろいろな側面があるのだよと。要するに、尺度というのは下位尺度構造というのがありますので、単一構造ではないのだと。いろいろな側面、さっき言った道具依存の側面と情動依存の側面、情動依存の中にもこういういろいろな構造があるのだよということを多次元化することによって、単純な話ではないということ逆を主張していくというのが一つはあり得るのかなと思います。

実際に使うときにも、今の尺度でもやはり尺度を作った先生方は、こういうことには使っていないけれどもこういうことには使わないほうがいいよという作った側の方のステートメントがありますので、使うときには必ずそれを確認して使うという本来心理尺度を使うときの作法といいますか、それをきっちり守らないで使うというのは学会指針とか倫理綱領とかでもっと強力にしていける必要があるのかなと考えております。

○小塚座長 ありがとうございます。特に唐沢委員からは政策面でも尺度に依存し過ぎると危険だという御指摘もいただきました。重要なことだと思います。

そのほかいかがでしょうか。まだお時間はありますので。

馬籠委員、お願いします。

○馬籠委員 ありがとうございます。非常に勉強になりました。興味深く聞かせていただきました。

AIを使わないという選択肢がない上で、AIの進化としては個人の信頼を得られるかがポイントになってくるかなと思っておりまして、その面で言うと、信頼が置けるAIとはというところを考えていくのが一つあるのかなと思っております。例えばこれからどんどんインスタントチェックアウトというような機能が出てきて、AIと話した上で購買を促進させるみたいなものというのは非常に出てくるだろうなと思っております。その際にこのAIは信頼を置けそうである、というポイントがあるとよいのかなと思っております。信頼できるAIとは、というところを啓蒙していく必要があるのかなと考えているのですが、信頼のポイントという点に対して、何かしらお考えになられるところはありますか。

○野村委員 AIに対する信頼感と人間に対する信頼感、他者に対する信頼感というのは実は結構相関しておりまして、そうすると、人間の他者への信頼は何で決まるかというときに、その規範というのですか。それをAIに対する信頼の規範にそのままある程度持ち込

めるのではないかなと考えております。

ベタな話なのですが、何でも肯定してくれる他者とか、何でもポジティブな意見ばかり言うてくる他者というのは、多分我々は信用できないと思うのです。そうすると、AIに関してもバランスよく情報を提供してくれる。こういう商品は何かあるかと聞いたときに、こういうのが代表的だけれども、これにはこういうメリット、デメリットがあるよというポジネガ両方出してくるようなものが信用されやすいのかなというふうな人間基準ですかね。人間の信頼基準はある程度参考にできれば、それに基づいてAIをつくれば信頼は高まるのではないかと。簡単な予想ですが。

○馬籠委員 ありがとうございます。非常に参考になりました。

○小塚座長 ありがとうございます。

そのほかいかがでしょうか。

岡崎委員、お願いします。

○岡崎委員 大変貴重なお話をありがとうございました。特に対話型AIへの依存を定量化するというのはすごく重要な方向性だと思います。

特定の人物への依存ですとか、あとは例えばゲームへの依存というのは、診断とか治療があつて病院に行ったりということがあるのだと思うのですが、対話型AIに関して診断とか治療がどのくらいできるというのか、考えられているのか。あと、例えば特定の人物への依存に関する診断や治療とどのくらい違うと考えたほうがいいのかとか、御存じでしたら教えていただきたいです。

○野村委員 それに関しては私は何の役にも立たない状態にございます。というのは、私自身は臨床経験を全く積んでおりません。訓練ありませんし、だから、そこに関しても全く何も申し上げることができないのですが、実はAIに関する診断と治療というのは当分かかるだろうなと。というのは、ゲーム依存でさえも、言ってみればまだ診断基準がDSMで確定していない状態です。その中でももちろん久里浜医療センターとかああいうところでネット依存の治療をされていますけれども、あれは正式な治療診断というのはどこまでやっているのかというのは、正直、中を見ていないので分かりません。ただ、問題は確かに起こっている。起こっているから何らかの治療メソッドは存在するというのはインターネットやゲームでは多分あると思います。ただ、正式な診断基準とかは遅れているので、ましてやネットでさえそれなので、AIというのはどうなるのかなと。

さらに怖いのは、そうこうしているうちに新たなAIがどんどん出てくる。そうなってくると、本当に依存という問題をますます把握しにくくなるのではないかなと。でも、問題は起こるからというので、これから臨床の方々は大変になるのではないかなというネガティブな予想を立てております。

○岡崎委員 ありがとうございます。

○小塚座長 加藤委員、お願いします。

○加藤委員 私、質問があるのですが、この依存というのは子供とか高齢者とか年

年齢層によって何か特徴的なものとかはあるのでしょうか。

○野村委員 ゲーム依存に関しては、若年層がなりやすいという統計は上がっているそうです。

○加藤委員 それはゲームを利用しているのが子供が多いとかではなく、高齢者がゲームをやってもあまり依存はせずということですか。

○野村委員 いえ、高齢者はゲームをあまりしないのだと思います。

○加藤委員 分かりました。

そうすると、例えば年齢的な要因で依存傾向が何か違う、年齢特性というか、そういうものはあるのでしょうか。

○野村委員 恐らく脆弱な若年層、JSTのシンポジウムでもありましたが、脆弱な人々というのは引っかけやすいというか、陥ってしまいやすいというのはありますので、やはり若年層へのリスクというのは上の層よりも高いかとは思いますが。

○加藤委員 ありがとうございます。

そうすると、今、私はたまたま年齢層ということを上げたのですがけれども、大切なのはどこに脆弱性があるかということをしちゃんと特定した上で、その脆弱性に対してどういった依存が発生し得るのかということを考えるのが重要でしょうか。

○野村委員 多分そうだと思います。ただ、そのためにはやはり依存に陥っている状態というのを判定する何らかの方法ですか。それは心理尺度でもいいですし、何らかの外的基準でもいいので、それを確立しないことには手も足も出ないなというところがあると思います。

○加藤委員 ありがとうございます。

○小塚座長 いろいろ重要な御指摘がありましたけれども、そのほかいかがでございましょうか。

何かありますか。坂下委員、どうぞ。

○坂下委員 どうも御発表ありがとうございました。とても理解できる内容でした。

行為依存と人間関係依存という話があるのですがけれども、行為依存で見たときに、生成AIの対話型AIというのは、いきなり敬語では会話しませんよね。やはり我々は普通の平常語を使って何かを入力し、そのプロンプトに対して平常語で返してくる。これというのはとても親密感があって入りやすいのだらうと思うのです。そこで多分行為依存が発生する。人間関係依存というのは、普通、私たちは検索エンジンを見ながら何かを検索したときに、その結果を納得するのに多分自分で考えていると思うのですよ。誤ったデータが出てきたり、全然当てどもないデータが出てきて悩むことも多いと思うのです。でも、生成AIというのはデータを全部つなぎ合わせてくれて文脈を作ってくれるから、多分納得はしやすい。その中で人間関係的な依存が発生してしまうような気がちょっと思います。尺度を考えると、そのようなインターフェースなどのことも考えながら、考慮に入れながら考えていかないといけないような気がちょっと思います。

以上です。

○野村委員 ありがとうございます。

実は尺度を今つくっている最中なのですが、ある方がつくられた人間関係依存尺度というのがありまして、そこの項目を参考に項目を生成して、取りあえずβバージョンまで何とか、項目生成まで何とか終わって、これから妥当性を検証しようとしているところです。確かに人間関係上の部分はとにかく甘えさせてほしいとか、面倒を見てほしいとか、仕事を代わりにやってほしいとか、そういうむちゃくちゃなものがあるように思われますので、そこをこれから検証していきたいと思います。

○小塚座長 ありがとうございます。

そのほか、あるいは2回目の御発言を御希望の方もいらっしゃいますか。

私からも1点お伺いしたいのですが、仮に生成AIに対する依存が問題であるとして、この問題はいわゆる基盤モデル、あるいはそれに近いほうで対処すべきことなのか、あるいはそちらで多く発生する問題なのか、それともアプリケーションで使っていく側面で多く発生する、あるいは対処すべき問題なのか考えたときに、どちらにフォーカスしていくべき問題なのでしょうか。

○野村委員 それは依存の一つの要因と考えられる迎合性がどのレベルで調整可能なのかということに尽きると思います。それはモデルの基盤部分でそれが確固たるものとしてあるのであれば、そこをいじる必要がありますし、それがあった上で上層のところで、サーフェスウェブのところでそれが調整可能であれば、それはそれで対処可能と思います。

○小塚座長 ということは、どちらもあり得る。むしろ使われ方あるいは依存の現れ方によるということ。ありがとうございます。

大塚委員からお手が挙がっています。いかがでしょうか。お願いします。

○大塚委員 ありがとうございます。

今の小塚座長の御質問とほぼ重複するところではあるのですが、スライドでいうと9ページ目におきまして、対話型AIへの依存の大きな部分においては、AIの迎合性が重要なポイントとして挙げられているかなと思います。そういたしますと、個人への依存のほうに対処するのではなくて、AIの迎合性のほうに対処するということもあり得るのかなと思います。要はあまり迎合的でないというか、こういったナルシズムに影響を与えないようなAIを作れとある種ガイドラインなりで要求していくということもあり得るのかなと思いました。

そういったときに、野村委員の御研究がどういうAIが望ましいのか、対話型AIを作るときにどういう点に注意すればいいのかということにも役に立つのかなと考えまして、対話型AIにどういう対応策といいますか、依存を発生させないような策を組み込むとよくなるのかということを考えていたのですが、この辺りについて何かお考えがあればお聞かせいただきたいなと思いました。

○野村委員 ちょっとネガティブな話で申し訳ないのですが、これも先週のJSTのシンポ

ジウムで語られていたのですが、企業が生成AIを出している以上、企業は迎合性を止めないということをおっしゃっておられました。だって商売ですから、使ってもらえないとどうしようもないので、それは譲れない線だと。

そうすると、こちらとしての対策は、その迎合性に化学反応を起こしてしまうようなパーソナリティーの方々はどういう方ですかと。パーソナリティーのリスクパターンというのですかね。それを探していくというのが当面の仕事になるのかなと考えております。

○大塚委員 ありがとうございます。

おっしゃるとおり、企業は積極的に止めようとは思わないと思うのですが、ある種法によってこういうふうな表示にせよと強制していくことはあり得ないのではないのかなと思います。そのときにどういうふうな強制が一番望ましいのか、要は企業への負担もあまり大きくないレベルで、しかし、個人の依存をなるべく押さえ込むような法のデザインをどういうふうに考えていけばいいのかという点が気になったというところでした。ありがとうございます。

○小塚座長 ありがとうございます。

恐らく私が先ほど質問したことと大塚委員が今言われたこととつながっているのだと思うのですが、基盤モデルのところでは仮に迎合性が、それは基盤モデルを開発する人は企業である以上、それを目指すのです、それで使ってほしいのですと言うのですが、今、野村委員がおっしゃったように、結局それが単に長時間使うということだけであれば、依存として対処すべき問題なのかどうかということは出てくると思うのです。ところが、それがどこかのアプリケーションに入ってきて商品の購入に誘導するとか、あるいは何か個人情報取得することに誘導するとかそういうことになってくると、むしろそこでそこそが問題だという考え方もあり得ますし、ただ、逆にそう誘導することこそが企業の本質なので、そこではなくてむしろ基盤モデルのところでの依存性あるいは迎合性に少しブレーキをかけてもらわないといけないという考え方も出てくるだろうと。恐らくそういうことだと思いますけれども、どちらにせよ、難しい問題があるということですかね。

○野村委員 日本だけで対処できる問題ではないと思います。

○小塚座長 御指摘のとおり、その問題はもう一つあり、しかも、サービスはグローバルでありながら、受け取る人間の側は文化によって違いがあるという先生がおっしゃった問題もあるということで、非常に困難な問題があらうかと思います。

今日は大変勉強になりまして、私どももよく分かりました。ありがとうございます。一旦次に進ませていただきまして、またそれも踏まえて議論したいと思います。

ということで、このAI技術と消費者問題について第2のプレゼンテーションということで、今日は坂下委員にもプレゼンテーションを御用意いただいております。同じぐらいの時間で御発表いただいて、また十分議論をしたいと思います。よろしくお願いします。

○坂下委員 本日は報告の機会をいただきましてありがとうございます。JIPDECの坂下でございます。

当協会はプライバシーマーク制度の運営以外にもパーソナルデータの調査研究などもやっております。例えば昨日も出生数が67.1万人という報道がありましたが、今、人口学では日本というのは35年で4割人が減っていくと推計されています。そのような中では、従来マスでやっていたサービスを個別にやらなくてはならないという課題があり、そういう（その時にプライバシーをどのように保護するか等）調査研究などをやっております。

それでは、早速報告に入りたいと思います。次のページをお願いします。

冒頭で内閣府のアンケートの報告もありましたが、当協会もAIの利活用の調査を先月実施しております、今日オンラインで御覧になっている方も新聞等で御覧になったかもしれませんが、多くで報道もされました。内閣府のアンケートですと10代、20代の若年層が生成AIを使って人生相談をしているという報道があったのですが、我々の調査ですと、人間とAIのどちらを信じますかという問いに対して、60代から70代の女性の47.8%はAIのほうがいいですという回答をしております。有識者によると、気軽さとか感情的なもつれがないので利用度が高いのではないかという評価をいただいておりますけれども、これらの調査を見ますと、全年代に対してAIは浸透していて、AIの過度な依存とか判断力への影響とか心理的な影響を見ていかなくてはならないということを物語っていると思っております。

次のページを御覧ください。

今回、この報告に当たりまして、事務局から2つテーマをいただいております。1つ目が事業者の経済的利益の最大化のためにプロファイリングの巧妙化と依存につけ込む不当な広告・勧誘が生じるケースはありませんかというのと、2つ目がAIエージェントが独立のサードパーティーによって提供される場合でも、実は事業者側から一定のインセンティブが出るケースはないですかというものです。

私たちは日頃こういう調査をやっている、事業者のヒアリングなどもやっておりますので、実際に該当するであろう事例を今日はお持ちしました。ただ、事業者名は秘匿する義務を負っているので出すことはできません。また、全ての事業者がこういうことをやっているわけではないということを冒頭に注意を促しておきたいと思います。

では、次のページをお願いします。

まず、経済的利益を最大化するための取組の事例ですが、前提として、企業が利益を最大化するのは当たり前のことです。今日御参加の委員の方々も、御覧になっている皆さんも、どこかの組織に所属している方がきっと多くて、その組織は利益を最大化しようとして、その中で私たちは給料をもらって生活をしています。ただ、その最大化する活動の中で、度が過ぎるとか、社会通念上これはいかんでしょうという問題がある場合が現実にあるということでございます。

左側のほうが健康食品の販売の例でございます。消費者保護と人格権とテーマに書きましたが、顧客の対象として例えば高齢者などで認知症を心配している人とか、家族に迷惑をかけたくない、元気でいたいという人とか、孤独な人という人たちに対して広告を

出したいと思います。収集される情報の例は、例えば検索履歴、病気関連をよく見ているとか、また、一人暮らしの推定です。昼間長時間見ているとか、子供との接触頻度、これは閲覧時間のインターバルを分析して、多分この人は親の代わりに保育園に行っているというようなものを見ている。そこからプロフィールを事業者は作ります。認知症を心配していて、独居であって、消化器系の病気を抱えている人だというのが仮に分かったといたしましょう。そこに対して不安を煽る広告とか、即決を煽る広告とか、御家族に迷惑をかけないためにもという前置きをした広告とかというものが出てきます。これは、同様なものは、当協会の調査ですと投資とかリフォームとか見守りサービス、終活サービスなどでも見られました。

右側が先ほど野村委員からの御報告にもあったネットゲームでございます。これは依存傾向の介入というテーマにしました。顧客の対象は、心理的な刺激によって課金を頻繁にしてくれる人などを探しています。収集される情報の例としては利用時間帯、例えば深夜利用が長いとか、連敗後の行動ですね。負けたか勝ったかは結果が分かりますから、負けた後に画面の遷移を見るわけです。そこで衝動課金の頻度、課金に遷移する速度がどれくらい速いか。支払金額の差、例えば給料日だろうということを推定して、その翌日はたくさんお金を使ってくれるというものを推定していきます。生成されるプロフィールの例としては、負けた直後に高額課金しやすい人は誰かみたいなもので見ていくわけです。そこで、広告のほうでは射幸性を煽る広告、今だけとか、限定とか、あと1回で逆転できるとか、損失回避というような広告を打つことになります。これは、類例は当協会の調査だとソーシャルゲームとか、スポーツベッティングというスポーツの賭けですね。また、仮想通貨の投機や、違法なオンラインカジノなどに観られます、実際に当協会でも調査のためのアクセスをして確認しました。

次のページを御覧ください。

次はダイエット広告になります。こちらは承認欲求の利用というテーマにしました。顧客の対象はコンプレックスが強い人、また、同調圧力に弱い人、FOMOというのは事業者の方から聞いて教えてもらったのですけれども、Fear Of Missing Outといってずっとつながっていないと取り残される不安を持つ人、そういうような人が対象になります。

収集される情報の例としては、SNSに依存していないか、これは利用時間を見ます。閲覧履歴はコスメとかSNSとか診断とかクリニックサイトをよく見ていないか。もう一つが、失恋直後に失恋動画を見たり、SNSで失恋をした経験を書いていた、悩み相談サイトにアクセスをしているとかというのを集めているというケースがありました。

生成されるプロフィールの例としては、例えば失恋直後のような場合だと合理的判断が難しくなっていますので、購買につながられるのではないかと想定されるようです。

採られる手法の例として、ここで対話型AIが出てきます。行動UIーシコファンシー、相手の言うことを絶対否定しないという手法ーを使って心理誘導をしていくというものです。

類例としては、ダイエットやSNSの投げ銭ですね。また、高額スクールとか恋愛商法とい

うものが挙げられます。

(資料上) 1つ色を変えましたが、ダークパターンというのがあります。危険性は、プロフィールと連携することによって個人ごとに最適化されたUIが出されるということになります。例えばUIの例だと、解約が分かりにくいとか、自動更新とか、偽の残数表示をするとかというのがありますが、生成AIを利用した場合ですと、長期の会話ができますから心理特性を推定できるようになります。そこで最も断りにくい人は誰かとか、最も断りにくい表示は何かということで広告を出していくというのがあります。

実際にヒアリングの中でまとめたものですが、こういうものが事業者の一部の中で起きているということになります。

では、次のページをお願いします。

2つ目のテーマとして、AIエージェントが独立のサードパーティーによって提供される場合に、実は事業者から一定のインセンティブが出るケースはあるかというものになります。こちらについては、Perplexityという会社がBuy with Proというサービスをやっておりますので、それを報告させていただきます。このPerplexityという会社はAIの検索エンジンサービスをやっている企業です。同社がBuy with Proというサービスを始めました。調査目的で実際にやってみましたが2,000円ぐらいの部屋の中のライトが買いたいというのを入力し、そこでクレジットカードも登録しておく、決済までやって物が届くというものです。そういうサービスをやっております。

一番いいものを選ぶのに、通常、私たちは1つのプラットフォームに対するECサイトで物を買いますが、このエージェントというのは全てのECサイトを横断して見ていきます。その中で一番適当だというものを買ってくるのです。

ここで大きな問題が出ます。従来の広告モデルが崩れるということです。そのためか、このBuy with Proが出た後に、各プラットフォームは似たようなAIエージェントを出してきました。

このようなものが社会に浸透していった場合に、企業が経済的利益を最大化するときにおけるリスクは何かというものを簡単にまとめました。

一つは、従来はサイトに来てもらって、広告を見てもらって、そこで収益を上げていたわけですが、今後はそこに来店する事業者もそのAIが自分のところの商品を推薦してくれないと困るわけです。現状やっているサービスは販売紹介手数料は取っていませんが、一方でデータ分析はやっています。このようなことが起きてくると、AI推薦の最適化市場というものができるのではないかと私たちは今考えております。具体的には、AIが比較検討のためのツールから購入代理をしてくれるツールに変わっていったときに、その推薦の順位とか、提携の条件とか、実際の購買行動が極めてそこに大きな影響が出てくるのではないかと考えております。

振り返ってみますと、これは検索エンジンが出てきたときの問題と全く同じです。あのときもウィッシュリスト等に推薦されたものが出てきましたが、あの順番はいいのかとい

う議論が当時あったと思います。それが今このAIエージェントでもまた出てきたということだと思います。

次のページを御覧ください。

私たちのほうで今後こういうものを考えていかななくてはいけないのではないかとということを私案ですけれどもまとめてみました。

一つは、やはり広告に変化が見られます。従来の広告というのは、ネット広告もそうですが、属性ベースです。過去の分析であった、ある類型ごとに一斉送信する等という広告だったと思います。ただ、AIが活用されるようになって、心理状態をベースに個別に分析ができるように変化してきております。そこは行動予測がより精緻に行えるとなっていると考えられるのではないのでしょうか。

また、従来の「データをどれだけ集めたか」に加えて、「そのデータによって人が動いたかどうか」という観点が加わってまいります。このように見ていきますと、行動を予測されるリスクが発展すれば、誘導されるリスク、さらにその個人がスコア化されるリスクについて考えなくてはいけなくなるのではないのでしょうか。

ヒアリング等を基にしておりますが、下側にEC、動画やSNS、医療や健康、人事、自動運転で予測内容の例と使われているデータの例をまとめてみました。例えば医療や健康が真ん中にございますが、鬱傾向とか心不全の悪化、転倒リスク、再入院するかどうかというのは歩数のデータや睡眠のデータ、電子カルテ、会話、スマートウォッチみたいなもののデータも分析対象になっていると聞いています。

次のページを御覧ください。

もう一つは、この対話型AIという非常にインターフェースが柔軟かで慮った処理をしているように見えるものが消費者の代理になるような可能性がこれから出てまいります。そうしますと、AIに委任をしやすくする、またはAIに推薦されやすくする、さらには、AIが理解しやすい商品データを渡すように事業者が努力を進めていくというふうになるのではないかと思います。AIによる推進順位や提携条件が実際の購買行動へ大きな影響を及ぼすということを先ほど申し上げましたが、これはひとつ大きな考えるべきリスクではないのでしょうか。

一方で、従来の検索における広告はスポンサーの広告の表示や広告ラベルがありましたけれども、これはAIには対応していません。AIが消費者の代理人として動くことが想定される段階では、それらに応じた措置が考えられなくてはいけないのではないかと考えております。

また、利用する消費者の不安とか衝動性・経済状態・孤独・消費動向などから、AIが学習して最も契約しやすいタイミングで推薦・勧誘するようになってくると、それは行動操作になってしまうのではないかと思います。

下側にイラストを書きましたが、1993年にインターネットが商用化されて、2008年に日本でiPhoneが発売され、2014年はFitbit等のウェアラブルが出てきました。2025年はAI元

年とされています。この中で取られてきたデータ、私たちが提供しているデータというのは、インターネットしかなかったときは固定された場所で使っていました。そこでは時間が取られ、閲覧履歴が取られていました。購買履歴が取られて、登録段階で趣味趣向が分かっていました。スマートフォンはその状態で歩き回るようになりました。移動履歴が取られて、日時や閲覧履歴が取られていきました。そこにウェアラブルが加わりました。ウェアラブルは心拍とか血圧とか歩数とか身長、体重が分かります。ここにいたくないと思ったときは歩く速度が早いかもしれません。そういうデータが取られています。そこにAIが入ってきて、心理状態が推測されるようになっていきます。

このような個人情報がいよいよ使われる時代になって、そこから消費者を守るためのルールはやはり考える必要があります。ただ、実際には、今、消費者の心理的影響とか感情的影響というのがまだ散発的な研究に頼っていますから、そこを総括的にまず網羅的に見て整理をするということが必要なのではないかと考えております。

非常に簡単ではございますが、私からの報告は以上でございます。

巻末にここの報告で使った用語の定義はつけておきましたので、ネットで御覧になっている方々で分からない方はそこを御覧ください。

以上でございます。ありがとうございました。

○小塚座長 ありがとうございます。

用語の定義もなかなか味わい深くて、ヒト：個人は特定していないが個を特定している状態とか、非常に重要な問題を御提起いただきまして、また、議論の素材として非常に有益かと思えます。

どなたからでも、また、どの点からでも御発言、御質問等いただけますでしょうか。

野村委員、お願いします。

○野村委員 坂下委員、非常に貴重な情報をありがとうございました。

特に最後のところで、バイタルデータを取るのにApple Watchとかああいうのが出てきてという話がありましたけれども、結局、今のところAIを駆使しているのはビッグテックですよ。

○坂下委員 そうです。

○野村委員 そうなってくると、結局そのビッグテックがさらにバイタルデータを取るためのデバイスを駆使してくるというビッグテック一人勝ちになってしまうという傾向がますます強くなるように思うのですが。

○坂下委員 おっしゃるとおりです。私はFitbitを使っていますが、Fitbitは先般アプリケーションがGoogleアプリに変わりました。このような例からも、ビッグテックが益々データを集めているというのは事実です。ただ、データを集めるのはいいのだけれども、その使い方はこちらから制御しなくてはいけないという部分があります。そこを制度でやるのか、業界の自主ルールに任せるのかということは議論をしなくては行けませんし、なぜそれが必要なのかということは私たちが提言を出していかななくては行けないだろうと

思っております。

○野村委員　そのところの足並み、ハードローかソフトローかということになると、当然欧州とアメリカとアジア、我々のところで足並みがそろわない可能性がすごく高いように思われるのですが。

○坂下委員　御指摘のとおりだと思います。各国で同意ができるかどうかというのはとても大きな問題で、本件に限ったことではありませんが、アメリカ、欧州、日本では考え方が大きく違っている部分がどうしてもあります。そこではふわっと合意をしておいて、法域の中ではきちんと律するみたいなことを考えていかなくてはいけないのではないかと思います。

○野村委員　ありがとうございます。

○小塚座長　ありがとうございます。

今おっしゃったことは非常に重要な問題だと思いますが、国際情勢というのはただでさえ難しい消費者問題の中でまたとりわけ難しい部分でして、後で時間がありましたら私も申し上げたいことはありますが、取りあえずお手が挙がっている方に進みたいと思います。大塚委員、お願いいたします。

○大塚委員　ありがとうございます。

坂下委員、実際の事例を基にした大変貴重な御報告をありがとうございました。特に事例1の経済的利益を最大化するための取組の例というところで、こういう実際の事例があるのだということで、我々が取り組まなければいけない消費者問題が具体化されたのかなと思います。

私からの質問は、その点で何を問題視すればいいのだろうかという点についてです。今回、3ページ目、4ページ目で実際に取り上げていただいた例は、こういったデータを基にこういう手法で購買意欲を誘って煽っているのだよというものでした。例えば3ページ目の例でいえば、高齢者の不安を煽ったり、あるいは射幸性を煽ったり、こういった手法だとそれは駄目でしょうと言いやすいといえますか、駄目な部分が見えやすいものなのかなと思います。

これを基に、7ページ目のスライドでは行動操作になってあまりよろしくないよねという御指摘をいただいたと思いますけれども、それはこういった事例ではそうかなと考えました。ただ、データの収集がより拡大していく、あるいはAIの利用がどんどん拡大していくとなりますと、ここに着目して購買意欲を誘っているのだというポイントが人間には見づらくなってくるのではないかと思います。要は、今御紹介いただいた例では、人間から見ても、我々から見ても、そこに着目しているのだというのが分かって、かつそれがあまりよくないねというのが分かるという例だと思うのですが、そうではなくて、いろいろなデータを分析して、こういうタイミングでこういう広告を打っているということなのだけでも、しかし、何でそれが購買力を煽るのに使われているのかが人間にとっては分からないというケースがどんどん出てくるのだと思います。それが射幸性を煽ると

か、そういった言葉に還元できないということになる。しかし、全体で見ると行動が操作されているとも見受けられるというときに、それが果たしてよくないのか、消費者問題であり、何か法によって対処すべきなのかということが今後考えていくべき課題になっていくのかなと思いました。

質問というのは、そういったデータ収集やAIの利用拡大に対応してどういうふうの問題を捉えていけばいいのかという点で、今回の御報告のさらに先の話となりますけれども、坂下委員に何かお考えがあればお聞かせいただければと思っております。よろしくお願いいたします。

○坂下委員 ありがとうございます。

最後のページに今投影されていますけれども、いろいろなデータが取られているようになっているわけですが、例えばウェアラブルのデータというのは、消費者側から見た場合にはここで自分の知らなかった自分の姿が分かることもあります。そういうことを知るということはまず大事だと思います。

また、事業者側は多分データをたくさん集めていって分析をしているのですが、どんなサービスであっても答責性は必要です。なぜそうになっているか等は求めていく必要があると思います。

時間をかけてやっていくものとしては、教育の問題があります。今、高校でも情報が始まりましたが、情報の中ではプライバシーマークも紹介されていて、個人情報の保護とかプライバシーの保護とかという議論も中に掲載されています。そうやって若い人たちの教育度を上げていくということも大事だと思います。

これは何かの法律をつくって全てが変わるということではなくて、様々なパズルを組み合わせると1つの絵に変えていかないと、社会全体は変わらないと思います。技術は発展してしまいますので、それをいかに社会実装していくのか、その中でいかに人間と融合させるかということが課題になるのだと思います。そのためには、こういう検討会という場は今後も続いていくことになるのではないかと私は思っております。

以上です。

○大塚委員 ありがとうございます。

AIの利用というかデータの利用をする場合の答責性を上げていくというのは必要なことだと思うのですが、AIはブラックボックスであると言われており、どこまで説明を可能にするのか、そして、何らかの説明が可能だとしても、その説明が果たして適切と言えるのかというのはさらに残っていく課題のような気はいたしまして、どういう方向性で答責性を上げていくということなのでしょう。

○坂下委員 答責性については、まだ国際標準の世界でも議論がされていますので、一概にこうだという回答は現状ではありません。ただ、実際に例えばトラストというのを経済学で考えたときには、その主体が大丈夫かということと、その所属集団が大丈夫かということと裏打ちする制度があるかという観点で見たりします。ですから、そのような知見を

援用しながら、補助線を引いて適切な制度を考えていくということになるのではないでしょうか。

○大塚委員 よく分かりました。ありがとうございました。

○小塚座長 ありがとうございました。

その次は丸山座長代理からお手が挙がっています。それから唐沢委員ですね。

ということで、丸山座長代理、お願いします。

○丸山座長代理 興味深い御報告をありがとうございました。

私からは、スライドでいいますと4ページの行動操作の問題と、後でAIエージェントの問題について質問させていただければと思います。

4ページのダイエットの例のところに対話型AIの利用というものが出てきております。ここでは不安を煽るというよりは迎合とか共感を示すといったことによって働きかけるといったことが行われているということだったのですけれども、結果としまして、結局契約などをした消費者、ユーザーが合理的意思決定していないであるとか、本人も後々後悔するといった意思決定における問題が惹起されていると理解してよいのかどうかというのが、質問の第1点になります。

第2点としましては、最後のまとめのほうのページになりますけれども、AIエージェントが用いられていく未来というところがあったと思います。この点について、具体例で挙がっていた例も含めて伺いたいのですが、本来消費者側が使うAIエージェントとなると、消費者本人の利益の最大化に機能すべきではないかとも思われるのですけれども、売り手、事業者の利益の尊重になっている、言わば利益相反になっているような現状というのは現実にあるのでしょうか。これが1点。

もう一つ、AIエージェントに関しては、消費者側のAIエージェントを利用すれば環境操作の問題とかはなくなるように思えるのですけれども、逆に消費者側が使うAIエージェントについて、何かAIエージェント特有の性質を狙った操作といったものが現在懸念されている状況というものはあるのでしょうかというこの2点はAIエージェントについて教えていただければと思います。よろしくお願いします。

○坂下委員 どうもありがとうございます。

最初の1点目の御質問で消費者問題に直結していないかということですが、実際にこれは相談がいろいろありまして、消費者問題として国民生活センターなどでもそういう問題が出ているということを伺っておりますし、当協会もプライバシーマークの苦情相談窓口があるので、本来は受けなっておりますが、そういう質問や相談が来ることもあるので、顕在化していると思っております。

2つ目のAIエージェントのほうですが、AIエージェント自身が何か消費者サービスの中で今使われているかという、まだ実際に使われているものはほとんどないと思います。ただし、今回AIエージェントの中で購買をやるというものが出てきましたから、ここからスコープしたときにAIエージェントを中心とした事業構成の枠組みが変わるのでは

ないかというのが今回の問題提起です。従来のネットの中の広告モデルというものが崩れていって、消費者の代理人のようなAIエージェントを使ったサービスにビジネスモデルが変わっていくのではないかと。そのときにこのような課題があるのではないかとということで提示をしたものになります。

また、感情操作につきましても同様でして、AIについては人事評価などもまだしっかりできませんし、感情操作まではしっかりできるものはないのだろーと思っております。ただ、私たちというのは、UIが非常にソフトであり、慮った処理をしてくれているので、それによって感情移入していくところがあって、移入をしたときに結果として操作されている部分は今の段階でもあるのではないかと私は考えています。

これでよろしいでしょうか。

○丸山座長代理 大丈夫です。

○小塚座長 ありがとうございます。

それでは、唐沢委員、お願いします。

○唐沢委員 どうも御発表ありがとうございました。先生の御発表から、私たちの心理的な状態、動機づけ自体が商業的に利用可能なターゲットになっているという点を、非常に説得力を持って理解させていただいたと思います。

その上で伺いたいのですけれども、脆弱な人。ターゲットとなりやすい人について、事例を見せていただき、そういう脆弱な人を守るためにという発想も重要だと理解したのですが、自分は脆弱ではないと思っている我々も、常にターゲットになると思います。というのは、私たちは一時的に脆弱になることがよくあるからです。操作のされ方についての先ほど御説明では、脆弱な人と思われるカテゴリーについては、一定期間の相互作用のもと、情報収集がなされ、ターゲッティングされている。ところが、そのようなカテゴリーに当てはまらないが一時的に脆弱になった場合、履歴がないから比較的安全なのか、もしくは、非常に短い相互作用でも我々は情報を取られていって、それに基づき狙われるという状況があり、その点も対処する必要があるのか、または、現実に対処されているのかについて、教えていただきたいと思いました。

○坂下委員 どうもありがとうございます。

特にAIを悪魔みたいなことで考えているわけではないのですが、今、投影されていますけれども、AI元年になってきて、AIというのは大変多くのデータを読み込んでいきますので、類推がしやすくなっている部分はあると思います。従来過去に何かあった人であっても、一部の履歴があればそこから類推するということはできるのだと思います。しかし、その類推が誤った類推であれば、そこは直してもらう必要があって、その透明性は高めなくてはいけないのだと思います。ですから、技術が発展していって、データはもしかしたら今日から取られるのかもしれませんが、多分推測プロファイルは作られていて、ただ、それを我々は知ることができないので、そこに対してリスクはあるのではないかと考えています。

また、先生が御指摘のように、人間というのは24時間生きていて、必ず脆弱なタイミングはありますから、そこにこういうものがつけ込んでくるリスクはあるのだということが今日お話しできれば、まずは良いのではないかと考えております。

○唐沢委員 もう一点よろしいですか。消費者の代理人という言葉は重要な概念だなと思ってしまして、先ほどの丸山委員の質問とも関連するのですけれども、私たち消費者の側としては、AIが自分に忠実な代理人であること、例えば、私はこういうものが欲しいとおもえば、いろいろ比較して、私の状態からこれが必要なのだと、私に忠実に判断してくれれば、ある意味問題はないようにも思えますし、そういうサービスがあるならば使うでしょうし、今でもいろいろなサイトを比較して情報だけ提供してくれるようなサイトはありますよね。そこで我々は、すでに操作されているのかもしれませんが、AIと消費者との関係の未来像において、AIが必ず使われていく中で、AIが仮に消費者に忠実であったらどうなのか。先ほど申し上げたような私の消費をよりうまくガイドしてくれる状態が目指せるのか、または、こういう視座は先生の御専門からどう見えるのか、かなり楽観的な考え方で、消費者に忠実という概念そのものが非常に困難なのか、その辺りを教えていただけませんかでしょうか。

○坂下委員 多くの事業者は消費者を尊重したビジネスをやっていると私は信じています。ただ、それが過度に行き過ぎたときに、一線を越えたときに今日のような課題が出てきてしまうのだらうと考えております。

AIエージェントが消費者に寄り添ったものになるかどうかというのは、やはりサービス事業者のポリシーとか理念などに即したものになってくると思いますし、私たちはそれを例えば株式市場等で評価していくことになるのだと思います。

一方で、AIエージェントが非常に優れたものになり、執事のようになって、常日頃からガイドをしてくれるようなことになっていくと、だんだん人間というのは批判的思考がなくなって、自分で考える機会をなくしていくのではないかと考えております。人間として、これは人間がやらなくてはいけない部分だということをちゃんと分かっておく必要があり、その上で道具と付き合うということをやっていないと、どんなデジタル技術も社会実装には至らないような気がいたします。

以上です。

○小塚座長 ありがとうございます。

なかなか深い議論になってまいりましたが、田中委員からお手が挙がっています。田中委員、お願いします。

○田中委員 どうも御説明いただきありがとうございました。特に今表示していただいているスライドで指摘されていることは大変重要な問題であると感じています。

AIエージェントが消費者の代理になる可能性がある未来というのは、現実にもう起こっているのかもしれませんが、必ず起こり得る未来の一つであると私も考えております。そこで考えているのが、合理的な意思決定というものの何が合理的な意思決定なのかの定義

点を外部化するということが起こっていくのではないかと。つまり、人間にとってはヒューリスティックで、合理的な意思決定を代理人に任せるような外部化をするということが起こり得るような気がしています。

そういったときに、質問は、結論から言うと、消費者が消費者問題として認識することへのハードルが上がっていったら、現在は消費者の方から直接相談が来る状態ですけれども、それがなくなっていくというようなハードルが起こり得るのかどうかということについて坂下委員の御意見を伺いたいと思います。

もう少し具体的に言うと、このスライドの1、最初の三角のところですよ。AIに引用されやすくなる、また、AIに推薦されやすくなるということは実際に起こり得ると思います。そうすると、現在のAIエージェントは既存の広告、伝統的な広告システムに基づいて推薦をするようなアルゴリズムを作っていると思うのですが、これが一般的になっていくと予想されるのは、AIエージェントがハックされるということが可能性としてはあると思います。つまり、ダークパターンというものも人間の認知に向けたダークではなくて、AIエージェントがだまされやすくなるようなダークパターンというものが進んでいくとすると、AIエージェント自体も何をもってその出力が最適であるかを説明できなくなるリスク、AIエージェントの事業者自体もなぜその出力を出したのかを説明することが難しくなるリスクというものがあるように思います。

つまり、説明可能AIというものの表裏一体の話だと思うのですが、そうすると、誰も説明できないとか、誰も因果を説明できないし、複雑なモデルで、式は書けるのだけれども、そこに誰も日常のラベル言語を貼ることができないということが起こり得ると、消費者も何が被害なのかを認識することができないし、事業者もそれを説明することができないようになっていくと、消費者問題として社会が認識することそのもののハードルが上がるというようなことが起こり得るのかどうかについて御意見を伺いたいと思います。

○坂下委員 どうもありがとうございます。

資料の6ページというのを投影できますでしょうか。

今の御質問というのはとても重要なところでして、AIというものについてサービスに社会実装、ビルトインをしようとしたときに、多分レベルがあると思うのですが。例えば自動運転のところを御覧いただいて、事故リスクとか疲労というものが書いてあると思うのですが、これをいろいろな情報を取って分析しているわけです。そのときに何でそういう判断をしたのかというのは、これは車のメーカーが説明できなくてはいけないわけです。答責性が求められるわけです。

一方で、ECのようなところで次に買いそうな商品とか購入履歴とか閲覧時間とかカートの中身とかとありますけれども、ここで答責をどこまでするかはレベル感があると思うのですよね。グループの中から判断をして、こうしていると判断したのか、あなたの過去の購買履歴から見たのかというようなことはグラデーションが出てくると思います。AIの使

い方というのが適合するサービスと合わせたときにレベルがあって、そのレベルによって答責性のレベルも変わってくるような気がします。

ただ、実際には今はまだ社会実装されているものは少ないですから、具体的にこうだということは私も言えないのですが、ただ、きっとグラデーションがあって、そこでの尺度というものが先ほど野村委員が御報告された尺度にも関係するのではないかと私は思っております。

以上です。

○田中委員 ありがとうございます。

○小塚座長 ありがとうございます。

そのほかいかがでしょうか。

馬籠委員、お願いします。

○馬籠委員 ありがとうございます。非常に勉強になりました。

お話を伺っていて思っていたのは、商品を出品する側でも気にしていかないといけないところがたくさん出てくるということです。今後AIを例えばOEM的に企業の中で使っていくといった場合に、買わせる一方だと問題になってくるのではないかと感じています。現状のマーケティングの中ですと、CVRとかそういうところがポイントになって、どれぐらい売れますかというところがポイントになってくるのですが、そこではない尺度、例えばAIのさっき言っていた信頼性みたいな尺度というところを、商品を出品する側のほうでもしっかり持つことが大事になってくるのかなと思っております。それを考えたときに、野村委員に伺った際には、人間として正しい行動ができそうかというところが信頼ポイントかなという話だったと思います。今後気にしていくべき尺度、CVR以外の尺度をつくったほうがやはりよくなるのでしょうか。

○坂下委員 ありがとうございます。

広告モデルに限って申し上げると、ビジネスモデル自身が変わってくると思います。冒頭申し上げたように、日本という国は35年で4割人が減っていく国ですから、マスで広告を打つのではなくて、個別広告を打つビジネスが変わってくると思うのです。そのときにはAIのようなツールはきっと不可欠で、みんなが使っていくことになると思うのです。ただ、その道具を使ったときに、どこに配慮するのかということは考えなくてはいけなくて、それを広告代理店の方々と共同して、ソフトローかどうかは分かりませんが、（ルールを）つくってもらって運用するということはあり得ると思います。

だから、今までの人口がどんどん増えていく中でのビジネスではないという中で、たまたまこのAIというツールが出てきていて使えそうだという話になっていて、それを社会実装できそうだという文脈になっています。社会実装するということは安心・安全ではなくてはいけませんから、その安心・安全の尺度をどうするかということになってくる。それが今日の私の報告だったり、野村委員の報告だったりするのではないかと思っております。

以上です。

○馬籠委員 ありがとうございます。非常に参考になりました。

○小塚座長 ありがとうございます。

そのほかいかがでしょうか。

では、私のほうからも、途中で出ていた話にも関係するのですけれども、AIを提供する事業者側がデータをたくさん集めて、それに基づいてプロファイリング、パーソナライズしていくということはありますが、このデータ自体を消費者の側が言わばセルフプロファイリングですね。こういうふうにしていきたい。そのために、自己認識するためにデータを集めるというのは消費者にとって有益でもあるわけで、そうやって集めたデータを自分で管理して自分でセルフプロファイリングする。こういうような将来というのは描く余地があるのか、それともそれは技術的に難しいとお考えか、その辺はどうでしょうか。

○坂下委員 ここはデジタルの神学論争になってしまいそうですけれども、自己情報コントロール権と言われているものだと思うのですが、自分で自分の情報を管理する。これは本当に理想的なことだと思います。過去に日本の政策の中でも情報銀行だとかいろいろな取組が行われてきました。ただ、いまだに社会実装されておられません。そこはやはり利害があったり、ビジネスとして技術上無理があったりというのがあるからだと思います。ただ、技術は進歩しているので、いつかはできるかもしれません。

類例になるのは、インドが今、個人の情報を全部1つの箇所に集積して、政府が公認したデータベースの中で扱うように変わりつつあります（India Stackなど）。そんなものができればそこを使っていくのかもしれませんが、ただ、そこに至るに当たっては、個人のプライバシーの問題は当然ありますし、また、どこまで明らかにしていけるかという課題もあるでしょうから、もう少し熟議が必要なのではないかと思います。

私からは以上です。

○小塚座長 ありがとうございます。そういう意味では恐らくセッティング次第ということもあるわけですね。

そのほかいかがでしょうか。2回目の御発言でも結構です。

では、お待ちしている間に、途中で私が申し上げたいと申ししていたことは、国際的な考え方の違いをどうするのかというのは、これは本当にAIの時代の、大問題でして、従来は国際的にハーモナイゼーションということがよく言われてきたわけですね。要するにルールを1つにしましょうと。完全に1つでなくても、ほぼほぼ同じようなものに収れんさせていきましょうと。最近、特にAI周りではそれがかなり悲観的になってきて、言われているのがインターオペラビリティということです。要するに、例えばヨーロッパはヨーロッパの中、アメリカならアメリカ、日本がどこに位置するか分かりませんが、それぞれの政策判断というのは違っていても仕方がないけれども、事業者の側はグローバルに展開したいので、せめてここを調整すればこの領域には入っていただけます、この国に来たときにはここを調整します、その接合点だけは明らかにしてほしい。このインターオペラビリティという概念が出てきていて、そういう意味でいうと、一步引いたレベルで、し

かし、何とか国際的にルールのコ存、調和が確保できないかということが言われていて、恐らくそうなっているかざるを得ないのではないか。そういうことになっているということを申し上げたかったということです。

いかがでしょうか。坂下委員にお聞きしたいこと、あるいはそれに触発されてこの調査会のテーマとの関係で御発言になりたいこと、御提案等がありましたら。

では、坂下委員、先にどうぞ。

○坂下委員 今の国際ルールの小塚座長がおっしゃったものというのは、私も国際標準をやっていてとても難しいところがあります。簡単に言うと、レシピを作って、そのレシピどおり作っているから大丈夫だというのが国際標準の考え方だと思うのですけれども、そういうものでやれるかというやはり限界がある。行為をマネジメントとしてISOでルール化をして展開してはどうか。これは一定程度うまくいっていると思います。また、日本の場合だとEUと日本は個人情報の保護法については充分性認定を受けていますから、同じレベルの法律であるということでデータのやり取りがやられている。やはりやり方はグラデーションがあります。そのグラデーション自身をどういう目的のときにはこの色がいいということを決めてビルトインしていくのが社会実装上は妥当なのではないかなというのが私が考えていることでございます。

以上です。

○小塚座長 ありがとうございます。おっしゃるとおりだと思います。そこをまさに見極めていくというのが非常に重要なことだと思いますが、大塚委員からお手が挙がっていますので、大塚委員、お願いします。

○大塚委員 ありがとうございます。

今の話と関連して、プラットフォームが国際化している中でどういうふうに法を実装していくのかという点について一言だけコメントを申し上げます。

先ほどの御報告の中で、最後に野村委員が悲観的なことをおっしゃっていたように記憶しておりまして、国際的な企業の行動に対応するためには法では難しいので、もう少し個人の側面から対応していくべきだというお話だったかと思います。それはそれでそのとおりな面もあるかなとは思いつつも、ただ、日本における法というのが全く対応しなくていいのかというと、そうではないような気がしているところです。もちろん日本の法によって対応することには限界もあると思いますが、だからといって手をこまねいてはいけないように思います。

また、国際標準というのが必要だということもおっしゃるとおりで、EUではこう、アメリカではこう、日本ではこうとばらけていると、企業に必要な以上の負担を強いてしまうということはそのとおりなのですけれども、それでも手放しでいいのかというと、それはそうではないだろうと思います。そのときに日本としてどういうふうに対応していくべきなのかということで、いろいろな考え方があるように思いまして、これまではEUやアメリカで議論が固まってきた段階で、それを研究して、遅ればせながら日本ではEUあるいはアメリ

かに倣ってこういうふうに規定を置いていこうといった姿勢を取っていたような気がいたしますけれども、やはりどんどん技術の進展が速くなっているということだとすると、それでは遅くなってしまうという可能性はあるような気がいたします。そうすると、数年後の国際標準には合わない可能性、そういったリスクをある程度負いつつも、もう少し積極的に規制というか法的にいろいろなことを考えていくということが必要になっていくのではないかなと考えております。

なので、個人で対応するとか、あるいは技術のほうに働きかけるといったことももちろん重要なのですが、それと並んで法についても検討していく必要があるように思います。そのときには、どういうふうに法をデザインすると個人の利益をちゃんと守れるようになるのか、逆に企業の側にとっては、企業の利益を必要以上に減少させないためにはどういうふうに法をデザインしていくべきなのかといったようなことを考えながら検討していく必要があるのかなと思います。

以上です。

○小塚座長 ありがとうございます。

どうぞ、まだまだお時間はありますので、御発言がありましたらお願いいたします。

では、加藤委員、お願いします。

○加藤委員 坂下委員、今日はありがとうございます。

AIエージェントの問題なのですが、これが近いうちに社会実装されるであろうなと実感しています。昨年、既に国連で国際消費者機構がAIエージェントの問題点を指摘していきまして、消費者問題の次元がかなり変わってくるだろうと考えています。

これから技術はどんどん進んでいくと思いますので、このAIエージェントの消費者問題というところも視野に入れた政策とか議論といったものができたらなとは思っています。本当に次元が違うレベルになってくる。このAIエージェントを消費者が手軽に利用できるようになってしまった後では手遅れになる可能性があるものがたくさんあるなと思っています。以上です。

○坂下委員 どうもありがとうございます。

私もそこは同感で、従来の通販しかなかった世界に、インターネットが出てきてECが始まった。このときも同じようなドラスチックな大転換があって、消費者問題も発展というかレベルアップしてしまったと思います。

今回のこのAIのエージェントというものは、AIが出てきたときよりもAIエージェントというものが出てきたときにはまた同じようにレベルアップが起きていて、そこでは論点もまた違うものになっているのだと思います。ただ、原理原則は昔と何も変わってなくて、周辺環境変数が変わっているだけだと私は思います。そこには、原理原則を忘れないようにして、消費者の安心・安全を確保するのだという原理原則を守ったままで、その道具をうまく使っていくときにはどういうルールを考えていけばいいのだろうと、どうやってやればみんなが安心・安全だと感じるのだろうということをこういう場で議論していかな

くてはいけないと思います。ですから、先ほど申し上げたように、こういう会議というのは今後も続くのではないかと思うわけでございます。

以上です。

○小塚座長 ありがとうございます。

今日はゆとりもありますので、オブザーバーの先生方からももし御発言がありましたらお願いしますとだけ申し上げて、それを待つ間というのもあって、今のところにもコメントさせていただくと、やはりAIエージェントという形でも消費者の側の行動、選択、意思決定自体のところにAIが入ってくるということで、それを提供する主体というものもあり、その主体というのがまた背後に利益相反があるかもしれないということがブラックボックスになっていたりするという状況の中で、おっしゃるような本来の消費者政策とは何のためにあったのかということを常に意識していくことというのは非常に大事だと思うのです。今日坂下委員のお話でインセンティブだということが出ていたのは、恐らくその辺りを意識されていたのではないかと思います。

その前の大塚委員のコメントにもありましたけれども、もちろん各国の整備状況を見ながら日本の法制度を考えるというのが日本の伝統的な在り方でもあったわけですが、やはりそれでよいのかというところはありますし、とりわけこの新しい問題というのは各国の対応自体がまだ決まっていない。そういう意味でいうと、むしろそこに流れをつくっていくことが必要な局面でもあるわけで、そこを臆せずどういう判断をしていくのかということを考えるというのは非常に重要なことだと思いますので、このような会議は永遠に続くかもしれませんが、この会議としてもそれなりのものを打ち出していくということができればと思います。

いかがでしょうか。オブザーバーの先生方で何か御発言とかコメントとかをお持ちの方はいらっしゃいませんか。

坂下委員から御質問だそうです。どうぞ。

○坂下委員 野村委員の発表ですごく気になっていることが1個だけありまして、対話型のAIが親密な相談相手になって、依存関係ができ、本人の批判的な思考力が低下していった、自分で考える機会を失っていったらと、その人は社会的に孤立するのではないかと思います。孤立していったときに、相手がいませんから、さらにAIに入っていくってしまえば、そこでさらに過剰に信頼していくという悪循環が生まれていくようなことを思ったのです。そのようなリスクについて、学術の研究などで何か議論はなされているのですか。

○野村委員 まだそこまでは行っておりません。そもそも依存状態というのは特定がまだ完全にはできていないのと、先ほど申し上げた心理尺度の問題もあります。

ただ、これはネット依存のときに覚えておられる方もいると思いますが、『ネトゲ廃人』という本が出ました。先に症例が報告されるというパターンがあるのではないのかなと。そういう意味では、多分ちょこちょこあると思います。アメリカなどでオーバードーズを

AIに推奨されてやったというのとか、自殺を止められなかったとか、ああいうちょこちょこした事例が出て、日本ではこの間のような巨人軍の監督の問題とか、ああいうマスコミレベルでちょこちょこ話が出てきて、また騒ぎ始める、それで研究が本格化するという同じ運命をたどるのではないかなという気がしております。

○坂下委員 ありがとうございます。よく分かりました。

○小塚座長 ありがとうございます。

特に御発言の御希望はないということですのでよろしいでしょうか。

今日は野村委員にも坂下委員にも大変重要な点についての御報告をいただきまして、専門調査会は前に進むことができたと思います。

やはりずっと議論してきて私も思いますのは、AI、とりわけ生成AI、対話型AIがある意味で使う人、消費者の側の利益になるというところがあって、そこが結構大きいのです。坂下委員のお話でも、例えば従来の広告モデルが大きく変わっていく。広告モデルというのはある意味で言うと非常にラフな形で、個人の選好にかかわらず、好みの違いにかかわらず、一律にばんと商品を打ち出す。しかも、それは一方的な情報提供であるわけですね。ということに対して、ある意味で言うと、個人の好み、選好をきちんと把握して、自分が何を必要としているかということに合わせた提案をもらうというのは、消費者としてはありがたいことでもあるわけです。という面があり、しかし、同時にそれはまた危険なことでもあるという二面性がある。その中で、仮決めでいいけれども、どこかにやはり線が引かれなくてはいけないのではないかなと私は思っているのですけれども、そういうことをフォーカスしていただいたと思います。

前からそのことは思っていたのですが、それを確信したのが今日の野村委員の御報告でして、結局、依存の定義自体がそういうことなのだと。もともとメリットのあることをしていたら、それがあある段階でむしろ制御困難な状態になってしまうということ、それ自体が依存なのだという御指摘をいただきまして、そうすると、問題の本質というのはやはりそこにあるのかということで、両面あるので、難しいと諦めてしまわずに、臆せずにそこで一つのポジションを取っていくということが必要なのではないかと思っているところでございます。

という辺りで本日はまとめさせていただければと思います。皆様、いろいろな御意見、御議論をいただきましてありがとうございました。

また、野村委員、坂下委員におかれましては、御報告いただきましてありがとうございました。

ちょっと時間的には早いのですが、本日の議事は以上とさせていただきまして、連絡事項がありましたら事務局からお願いいたします。

《3. 閉会》

○江口企画官 本日はありがとうございました。

次回の本専門調査会につきましては、確定次第、事務局より御連絡させていただきます。

以上です。

○小塚座長 ありがとうございました。

なお、御指摘いただきました海外の動向とか、あるいはAIエージェントとか、今後この専門調査会でどういうふうに取り上げていくか、また事務局とも御相談させていただきたいと思いますので、ぜひよろしく願いいたします。

それでは、本日はこれにて閉会とさせていただきます。ありがとうございました。

以 上