

プラットフォームサービスに関する研究会（第10回）

1 日時 令和元年6月27日（木）10:00～12:00

2 場所 総務省第1特別会議室（8階）

3 出席者

（1）構成員

宍戸座長、新保座長代理、生貝構成員、大谷構成員、木村構成員、寺田構成員、松村構成員、宮内構成員、森構成員、山口構成員

（2）ゲストスピーカー

九州大学法学研究院・法学部 成原准教授

東北大学大学院情報科学研究科 乾教授

名古屋大学大学院情報学研究科 笹原講師

インターネットメディア協会 瀬尾代表理事

（3）総務省

谷脇総合通信基盤局長、竹内サイバーセキュリティ統括官、秋本電気通信事業部長、泉国際戦略審議官、竹村総合通信基盤局総務課長、山碕事業政策課長、大村料金サービス課長、梅村消費者行政第一課長、山路データ通信課長、中溝消費者行政第二課長、赤阪サイバーセキュリティ統括官付参事官、大内事業政策課調査官、岡本消費者行政第二課企画官

（4）事務局

三菱総合研究所 西角主席研究員、安江チーフコンサルタント

4 議事

（1）Facebookヒアリング（フェイクニュースに関する取組について）[非公開]

（2）米国におけるフェイクニュース対策の動向と議論（九州大学 成原准教授）

（3）ネットの誤情報に対する自然言語処理からのアプローチ（東北大学 乾教授）

（4）フェイクニュース拡散の仕組み 計算社会科学の見地から（名古屋大学 笹原講師）

（5）インターネットメディア協会の「信頼と創造」への取り組み（インターネットメディ

ア協会 瀬尾代表理事)

(6) 意見交換

【宋戸座長】 本日は、皆様お忙しい中、お集まりをいただきまして、ありがとうございます。定刻となりましたので、プラットフォームサービスに関する研究会第10回会合を開催させていただきます。

本日は、実は前半、後半の2部構成になっておりまして、前半は非公開として先ほど実施させていただいたところでございます。したがって、ここからは後半部分、議題で言いますと（2）以降を行います。

冒頭、カメラ撮りがございますので、少々お待ちください。

【岡本消費者行政第二課企画官】 よろしいでしょうか。お願いいたします。

【宋戸座長】 それでは、議事に入ります。今回は、前回に引き続き、オンライン上のフェイクニュースや、偽情報への対応について集中的に議論したいと考えております。そこで本日は、4人の有識者の方においでいただきまして、それぞれご発表をいただきたいと思います。

まず第1に、九州大学の成原先生から、米国におけるフェイクニュース対策の動向と議論。次に、東北大学の乾先生から、ネット上の誤情報に対する自然言語処理からのアプローチ。そして第3に、名古屋大学の笹原先生から、フェイクニュース拡散の仕組み、計算社会科学の見地から。そして、最後にインターネットメディア協会の瀬尾代表理事から、インターネットメディア協会の「信頼と創造」への取り組みについて、それぞれご発表をお願いいたします。そして、委員からのご質問等は、最後にまとめてお受けするということとさせていただきます。

また、本日は議事が多いため、それぞれ15分の発表時間をお守りいただきたいというふうにお願い申し上げます。そのため、学会的ではございますけれども、タイムキーパーをご用意させていただいておりまして、終了5分前、それから、終了の合図と、それぞれ鳴りますので、この点、報告者の皆様におかれましては、ご了承をお願いいたします。

それでは、まず資料1「米国におけるフェイクニュース対策の動向と議論」について、成原先生、ご発表をお願いいたします。

【成原准教授】 九州大学の成原と申します。情報法を専攻しております。本日は貴重な機会をいただき、ありがとうございました。時間が限られておりますので、早速、2ページをご覧ください。

私の発表では、米国における2016年大統領選挙から今日に至るまでのフェイクニュースと、その対策に関する動向と議論を概観することにより、我が国のフェイクニュース対策の在り方について、若干の問題提起を行いたいと思っております。

次に3ページをご覧ください。まずフェイクニュースの定義ですけれども、この定義については皆様御存じのとおり、必ずしも一致したコンセンサスがあるわけではないんですけれども、米国の文献などを読んでおきますと、その定義の要素といたしまして、情報の内容が虚偽であることに加え、虚偽であると知りながら公衆を欺くために意図して公表・拡散された情報であることが挙げられることが多いように思われます。従来のマスメディアにおいても、誤報というものは存在してきたわけですが、虚偽の情報が意図的に公表・拡散される点が、従来のメディアの誤報とフェイクニュースとの相違点とすることができるのではないかと思います。

もともと、この後、笹原先生のご報告でもご紹介があるかと思いますが、フェイクニュースにもさまざまなものが定義によっては含まれますので、必ずしもこの定義に限られるものではないんですけれども、1つの有力な定義として参照に値するのではないかと思います。

また、つけ加えますと、今日、フェイクニュースの問題というのが世界的に関心を集めているわけですが、メディア史の研究者らが指摘しているとおり、20世紀の前半から新聞やラジオなどを通じて虚偽のニュース、真偽の曖昧な流言、世論の操作を狙ったプロパガンダなどが流布し、各国の世論に影響を与えてきたということも留意する必要があるかと思えます。

次に、フェイクニュースがもたらす問題についてですが、こちらも米国の文献などを見ておきますと、例えば選挙の候補者に関する不正確な情報を流布するなどして、有権者の理性的な判断を妨げることにより、民主政治をゆがめたり、政治的分断を深めるおそれ。また、メディア等の発信する情報への信頼が失われるおそれ。さらに、外国政府が誤った情報を自国民に流布することで、民主主義と安全保障が毀損されるおそれなどが挙げられることが多いように思われます。もともと、これから米国の実証的研究を幾つか紹介する中에서도見ていくとおり、インターネット上のフェイクニュースの影響だけを過大評価すべきではないように思われます。現在の先進各国におけるポピュリズムや政治的分断には、複合的で構造的な要因があるからです。

4ページをご覧ください。次に2016年大統領選挙におけるフェイクニュースの影響について見ていきたいと思えます。これは皆様御存じのとおり、トランプ候補やクリントン候補に関する虚偽の情報が、大統領選挙の選挙期間中に多数拡散されております。このようなフェイクニュースが選挙結果に一定の影響を与えたとの見方もあります。

この2016年の大統領選挙におけるフェイクニュースの影響については、アメリカでは

多数の実証的な調査研究が行われております。ご専門の笹原先生、乾先生がいらっしゃるところで門外漢の私がご紹介するのも恐縮ですが、幾つか、2つほど有力なものを挙げておきますと、まず、こちらはニューヨーク大学とスタンフォード大学の経済学者らによる2017年の研究なんですけれども、彼らの研究によりますと、選挙期間中に米国の平均的な成人は1件から数件程度のフェイクニュースに接しており、クリントン候補を支持するフェイクニュースよりも、トランプ候補を支持するフェイクニュースに多くさらされているが、フェイクニュースによる投票行動への影響は、選挙結果を左右するほどではなかった可能性が高いと推測されております。また、興味深い点といたしまして、ソーシャルメディアは有権者の政治ニュースの主要な情報源ではなかったとの調査結果も示されています。

スライドの左側に添付している図をご覧ください。こちらは2016年の大統領選挙における米国の有権者の主要な政治ニュースの情報源について調査したものですけれども、意外なことにソーシャルメディアは13.8%でして、ケーブルテレビが23.5%、ラジオが6.2%、プリントメディアが8%、ローカルテレビが14.5%、ウェブサイトが14.8%、ネットワークテレビが19.2%ということで、テレビをはじめとする放送の影響力というのはアメリカでもまだまだ高いのだなということが、このグラフからも伺えるのではないかと思います。

次に、もう一つ紹介させていただきますと、こちらはハーバード大学の法学者や社会学者らの研究です。5ページをご覧ください。彼らはどういうアプローチをとったかと申しますと、ソーシャルメディアや放送などにより拡散されたニュース記事の関連性、つまり記事相互のリンクの関係ですとか、Twitterなどソーシャルメディアで一緒に共有されていると、そういった関連性を分析しました。その結果、どのような知見が導かれたかと申しますと、一般には2016年の大統領選挙における偽情報とプロパガンダの拡散の主要な要因として、ソーシャルメディア上のアルゴリズムに基づくニュース提供であるとか、AIを用いたビッグデータ分析に基づくターゲティング広告、ロシアなど外国から発信されたフェイクニュースといったものが注目を集めているんですけれども、彼らの研究によれば、それらは副次的な要因であって、根本的な要因は、過去30年ほどにわたる米国の党派及びメディアの分断化にあるという知見が導かれています。こちら5ページの左側にある図をご覧ください。いんですけれども、これは移民問題についてTwitter上で拡散されたニュースの相互関係を調べたものなんですけれども、青がリベラル・左派で、赤が保守・右派のもので、きれいに分極化しています。特にBenkler教授らによると、米国では右派のほうが先鋭化し、極端になっていて、この図でもBreitbartという右派のニュースサイト、極右のニュースサイトと

言われることもありますけども、そうしたメディアが情報のハブの中心になっていると。こうした長年にわたって築かれてきた分極化が、プロパガンダや偽情報の拡散の主要な要因になっているというのが彼らの分析で、参考に値するのではないかと思います。

もちろん、これ以外にもさまざまな調査研究があるかと思いますが、2016年の大統領選挙における世論形成において、ソーシャルメディア上の、あるいはプラットフォーム上のフェイクニュースが世論形成において主要な役割を果たしたという見方には、懐疑的な研究も有力だということは指摘しておきたいと思います。

次に6ページをご覧ください。2016年大統領選挙にロシアが米国の有権者向けにフェイクニュースを拡散するなどして干渉した疑惑が浮上し、捜査機関・情報機関より調査が行われてきました。本年3月、司法省のミュラー特別検察官が報告書を提出し、4月に一般に公表されたわけですが、こちらの報告書では、多数の証拠に基づいてロシアがFacebookなどソーシャルメディア上でキャンペーンを展開するなどして、大統領選挙に広範かつ体系的に干渉したことが明らかにされております。

また、連邦議会でもFacebook、Google、Twitterの幹部らを証人として召喚して公聴会が行われるなど、フェイクニュースへの対策の在り方などについて調査と議論が行われてきました。このようにアメリカでもフェイクニュースに対する調査、分析が政府においても進められているわけですが、アメリカでは表現の自由への配慮もあり、政府による規制には基本的に慎重なスタンスがとられております。

また、トランプ大統領は、みずからに対し批判的な姿勢をとるCNNなどの報道をフェイクニュースと呼ぶ一方で、Facebookなどプラットフォーム事業者による極右のアカウントの停止などについて、保守派の表現の自由への検閲だとして批判しております。

次に7ページをご覧ください。米国では合衆国憲法修正1条により表現の自由が手厚く保障されてきました。そのような伝統もあり、フェイクニュースへの規制には基本的に慎重な姿勢がとられています。連邦最高裁の判例によれば、修正1条のもとでは誤った思想（アイデア）のごときものは存在しないとされてきました。一方、事実については、事実に関する誤った言明は憲法上の保護に値しないとされてきましたが、それは自由な討論に不可欠なものであるとされてきました。つまり、真実の言明が萎縮することなく十分に保護されるためには、一定の範囲で虚偽の言明も保護する必要があると考えられてきたのです。

さらに2012年のAlvarez判決においては、虚偽の言明であるという理由だけでカテゴリーカルに、虚偽の言明が全体として修正1条の保護が否定されることはないという立場が明確にされております。

次に8ページをご覧ください。このように表現の自由が手厚く保障されてきたこともございまして、名誉毀損についても表現の自由との調整が行われてきました。米国の判例では、1964年のSullivan判決以来、自由な討論において誤った言明は避けられないとの認識を踏まえ、公共的な論点に関する討論が抑制されることのないよう、公人に対する名誉毀損について、民事訴訟において公人である原告が損害賠償請求を認められるためには、被告の現実の悪意を立証する必要があるという「現実の悪意の法理」が採用されてきました。すなわち、公人に対する名誉毀損については、たとえ虚偽の内容であったとしても、被告が虚偽であると知っていたか、または虚偽であるか否か無謀にも気にかけなかったことが証明されない限り、免責されることとなっております。もっとも、今日では、「現実の悪意の法理」の見直し論も出ております。

こちらは直接フェイクニュース問題とは関係ないんですけども、最近のとある名誉毀損事件の連邦最高裁の決定において、保守派のトーマス判事が、Sullivan判決は憲法の原意に基づかない政策判断だとして見直し論を主張しています。

また、これに関連してリベラル派として知られる憲法学者のキャス・サンステーンも、これはおそらくフェイクニュース問題を念頭に、うそが瞬時に広がるようになっている中、うそにより政治家の評価が損なわれ、民主主義が機能不全に陥るリスクが高まっているとして、同判決の見直し論に一定の理解を示しております。

次に9ページをご覧ください。米国の判例では、企業による選挙広告を含む選挙運動の自由が修正1条に基づき広く保障されてきました。一方、今日ではプラットフォーム事業者等が利用者に提供する情報の操作により、有権者の投票行動を誘導するデジタル・ゲリマンダリングも問題視されています。この代表的な事例としては、皆様もよく御存じの大統領選挙の際のFacebookからCambridge Analyticaに個人情報が出流して、ターゲティング広告などに使われたという疑惑を挙げるができるかと思います。

こうした問題を踏まえ、超党派の議員により「誠実広告法案」が提出されております。この法案は、月間5,000万名以上の米国民の閲覧者の訪れるプラットフォーム事業者に、当該プラットフォーム上で年間500ドル以上の広告費を支払った個人・団体による政治広告に関する情報、例えば広告の内容、ターゲット、閲覧数、広告料、広告主などの記録・開示を義務づけるとともに、外国の個人・組織による米国の有権者に影響を与える目的の政治広告の購入を防止するための合理的措置をとることを求めています。この法案については、Facebookらプラットフォーム事業者が支持を表明しております。

また、この法案は、基本的には民主党の議員が中心に支持しているんですけども、共和

党のマケイン上院議員も支持をしております。ただ、共和党の中には選挙運動、選挙広告を規制することについて慎重論、反対論も強く、この法案の成立の見通しは今のところ立っていないようです。

関連して、州のレベルでは、インターネット上の政治広告の透明性を向上するための州法などが制定されておりますが、違憲訴訟なども提起されております。時間がないのでこのあたりは割愛させていただきます。

次に11ページ、フェイクニュースと媒介者責任ということで、アメリカでは通信品位法230条に基づき、ユーザを含む第三者により発信された情報について、媒介者ないしプラットフォーム事業者に広範な免責が認められてきました。他方で、同条項のもとでプラットフォーム事業者は、第三者により発信された情報を編集・削除したとしても免責されてきました。この規定については、プラットフォーム事業者がユーザの投稿によって他人の権利が侵害されていることを知りながら削除しなかったとしても免責されるということで、あまりにもプラットフォーム事業者の責任を免除し過ぎではないかという批判もあるんですけども、反面でこの規定がプラットフォーム事業者による自主規制のインセンティブになってきたという評価もあります。

一方、2018年には性的人身取引を促す情報について、プラットフォーム事業者の免責範囲を限定する立法が制定されるなど、最近ではプラットフォーム事業者の責任を強化しようとする動向も強まっております。

12ページです。また、本年6月に共和党のHawley上院議員により、通信品位法230条の改正案が提出されております。この内容というのは、一定規模以上のプラットフォーム事業者が免責を受けるための要件として、他者の発信した情報を政治的に偏向した仕方で調整(moderate)していないことについて、FTCの認証を受けることを求めるものです。この法案は、トランプ大統領をはじめとする共和党の一部の政治家によるソーシャルメディア、プラットフォーム事業者に対する疑念、つまり、彼らは自分たちに対して政治的なバイアスのある削除方針をとっているんじゃないかというものが背景になっていると思われます。ただ、これについては民主党の議員から批判があるのはもちろんのこと、利用者団体や業界団体からも批判が寄せられております。つまり、政治的なバイアスの判断を政府機関に委ねることで、表現の自由が抑制されるおそれがあるといったことが指摘されております。

また、一般論と申しましても、アメリカでは今日でも政府がプラットフォーム事業者にフェイクニュース対策など、違法有害情報の削除を求めることに伴う過剰で不透明な自己検閲のおそれを懸念する議論も有力です。

次に、13ページをご覧ください。こちらについてはヒアリングなどでお聞きになっているかと思うので、省略します。

14ページをご覧ください。最近ディープフェイクが話題になっておりますけれども、ディープフェイクとは、AIを用いて作成される、主に映像によるフェイクニュースです。5月にはナンシー・ペロシ下院議長の演説を合成・改変した動画がソーシャルメディアで拡散されております。また、Facebookの創業者のザッカーバーグのディープフェイクなども拡散しています。ただ、ペロシ議長の「ディープフェイク」については、最初はディープフェイクだと報道されていたんですけども、現在ではこれは必ずしもAIを使ったものではなくて、動画の再生の速度を落としたに過ぎないということが明らかになっております。ということで、一部の論者はこれをチープフェイク、つまり、安上がりでできるフェイクだというふうに呼んでおります。

このことからわかるように、必ずしもAIを使っているということが根本的な問題なのではなくて、人間は簡単な技術で作られたものであったとしても動画や写真によって欺かれやすいということを認識しておくことが求められるのかと思います。

関連して、連邦議会でも公聴会が行われておりまして、来る2020年の大統領選挙に向けてディープフェイクが脅威となる可能性が指摘され、対策の必要性が説かれております。

15ページをご覧ください。フェイクニュースと放送ということで、アメリカでは公平原則が1987年に廃止されるなど、放送の規制緩和が進められてきました。これによって放送において論争的な見解、例えばアメリカではラジオの政治トークショーなどが盛んになったという評価もあるんですけども、他方で先ほど見たようなアメリカにおける政治的分極化を促す一因になってきたのではないかと指摘もあります。先ほども見たように、今日の米国でも放送は依然として大きな影響力を有しており、2016年の大統領選挙の世論形成においても、大きな影響力を果たしたことが明らかになっています。

最後に、我が国への示唆を短く申し上げます。まず、フェイクニュース対策の目的を明確化、具体化する必要があります。アメリカの議論や動向を見ておりますと、虚偽の情報だからということで一般的に問題にするのではなくて、それが例えば民主政治にもたらす害悪であるとか、具体的な害悪に着目して各論的にアプローチがとられているように思われます。このようなアプローチは我が国においても参考になると思います。

次にEvidence-basedの検討ということで、これは一般に今日の政策形成で求められていることですが、アメリカでは政府や民間の研究者らによって、フェイクニュースの影響について実証的な研究が行われております。我が国についても、このようなエビデンスを収

集した上で、対策を検討することが求められると思います。

また、先ほども見たようにインターネットやプラットフォームの本場であるアメリカですらも放送が依然として有力な、大きな影響力を持っていることを踏まえると、日本でも少なからず影響を持っていることがおそらく推測されます。ですので、インターネットのみならず、放送も含めてメディア横断的なフェイクニュース対策の在り方を検討していく必要があると思われます。また、インターネットに限って申し上げても、プラットフォーム事業者の責任を考えることはもちろん重要なんですけれども、一義的な責任はやはり発信者にあるわけですし、発信者への責任追及、また、我々インターネット利用者のリテラシー向上、さらに、広告主の自主的取組についても検討する必要があると思います。

それから、現行法を活用したフェイクニュース対策というのも重要になってくるかと思います。非常に簡単に申し上げると、アメリカに比べると日本では現行法により対応できる範囲が広いのではないかと思います。また、思想の自由市場の機能を維持・強化する取組として、ファクトチェックの支援などが重要になってくるかと思います。

最後に、プライバシー・個人情報保護法制との連携ということで、今日では個人情報保護委員会で個人情報保護法の3年ごとの見直しに向けた議論が行われており、有識者の先生方からプロファイリング規制の是非についても論点に挙がっていると伺っておりますけれども、先ほどCambridge Analytica事件に則して見たように、プライバシー・個人情報を適正に保護すること、あるいはプロファイリングを規制することは、間接的にフェイクニュースの拡散を防止して、健全な民主主義を維持することにもつながるということにも留意する必要があるかと思います。

駆け足になってしまいましたが、私からは以上です。

【宋戸座長】 高速でディープなご報告をありがとうございました。

次に、資料2でございますけれども、「ネットの誤情報に対する自然言語処理からのアプローチ」につきまして、乾先生からご報告をお願いいたします。

【乾教授】 東北大学の乾と申します。私の専門は自然言語処理と呼ばれるもので、情報科学、人工知能の中の一研究分野であります。その立場から、今日は少しネット上の誤情報に対するアプローチということでお話しさせていただきます。

2ページ目は私の紹介なんですけれども、それは飛ばしていただきまして、3ページ目ですけれども、誰でも情報を発信できるようなネット社会がどんどん浸透してからもう久しいわけで、それと同時に情報の信頼性をどう担保するか、誤情報の問題が顕著になってきております。それは今回のようなフェイクニュースが問題化する以前から問題でありまして、そ

うした問題に対して自然言語処理の立場から我々の研究分野、グループで取り組んでまいりましたことを、ここで幾つかご紹介させていただいています。

最初の左端は、例えばここではコラーゲンが肌にほんとうにいいかどうかみたいなことをネットに聞いて、そこからそれに対する賛否を整理して情報を集めてくるというような技術。あるいは東日本大震災が起りまして、その直後にTwitter上にデマがたくさん出回りましたので、そのデマを自動的に収集して、そのデマが拡散したり、あるいは収束したりする様子というのを技術的に解析するというようなこと。あるいはこの後、笹原さんのお話なんかでもあると思いますけれども、チェンバー効果のようなものがもう既に震災後の風評被害のようなものの中でもいろいろ観察されていたと。こうしたことを技術の立場からやっていたわけなんですけれども、そうした中でファクトチェック・イニシアティブという、前回の研究会でお話しされた楊井さんが事務局長で務めておられるファクトチェック団体を支援するNPOに、私も技術の立場から協力させていただいているような次第です。今日は、その活動を少し技術的な部分に焦点を当てましてお話しさせていただければと思っております。

次のページの青っぽいページですけれども、一言でファクトチェックと言いましてもいろいろな工程がありまして、ここでお示ししているのは先ほどご紹介しました、楊井さんが組織しておられるG o H o oというファクトチェック団体でやっているファクトチェックのプロセスでございますが、基本的にはネットから誤報のある種の疑義、これは誤報なんじゃないかというような、そういう疑義の情報を集めてきて、予備調査、本調査と上のほうにエスカレーションしていくというようなことをやっておられます。このプロセスを今、基本的には人間が全部やっているわけなんですけれども、AI技術のようなもので一部でもサポートしていければ、もう少し広範な調査ができるということで、FIJと一緒にそういうことをやっている。

もう少し具体的に言いますと、次のページですが、例えばあるニュースについて、右側の3つ目のツイートですけれども、他紙と違うが、両方裏づけてしているのかみたいに、誰かがSNSでおかしいと言っているような場合、この情報はファクトチェックにかける必要がもしかしたらあるかもしれないということです。そういう疑義の情報をG o H o oさんは端緒情報と呼んでいるんですけれども、まずはこうした端緒情報をネットから集めて予備調査、本調査へとエスカレートさせていくと。ところが、この端緒情報を集めるといっても、ネットの海の中からそうした情報を集めてくるのはなかなか大変なことでありまして、例えばデマとか誤報とかで検索しても、有用な情報というのはほんとうにわら山から針を探すようなものでして、ここにあるようにデマというキーワードがあっても、ほとんど疑義を呈してい

るようなものにはなっていないと。

8 ページはちょっとよくわからないと思いますけれども、実際にG o H o o さんではこうした作業画面で日々膨大な時間を使って、わら山から針を探す仕事をされておられました。そこをまずはF I J では支援しましょうということで、次のページの三角の絵ですけれども、F I J とSmartNewsさんと我々の研究室で組みまして、連携して、その部分の支援から取り組んでいくというようなことを始めているところでございます。

基本的なアイデアは次のページなんですけれども、疑義言説の候補を機械で絞り込むというところですが、ニュースがいろいろたくさんあります。左側にニュースがあります。それぞれのニュースについて、ツイートでいろんな人がいろんなことを言っているわけですね。そのいろんな人がいろんなことを言っている、その人々の反応を見て、疑義を呈している可能性が高いか、ファクトチェックにかけねばならない可能性があるかということが機械で見積もる。その可能性の高いものからランキングしまして、そのランクの上位のものを人間が実際にファクトチェックにかけたほうがいいかということを確認していくというようなことをやる。このランキングを自動的にやりましょうというのが我々の技術の話です。

次のページで、基本的に、細かいことはどうでもいいんですけれども、これまでにファクトチェックすべき記事と、ファクトチェック不要であるというふうに人が判断した記事、それぞれがそれなりの数たまっていますと、機械学習と言うんですけれども、それをお手本にしまして、次のページで真ん中に機械学習とありますけれども、人が○×をつけたものをお手本にして、どういうものが○になる可能性が高いかということを経験学習で学習します。学習が大体できてきましたら、下のほうから新しい記事が上がってくるわけなんですけれども、その新しい記事をどれぐらい○の可能性はあるか、どれぐらい×の可能性はあるか、どれぐらいチェックすべき記事である可能性が高いかということを経験学習で判定しまして、ランキングしまして、上のほうを人間が見てほんとうにチェックすべきかどうかということを経験学習するというようなことを日々やるという仕組みをSmartNewsさん、F I J さんとつくりまして、今それが稼働しているような状況であります。

13 ページは、人がそれを作業するイメージなんですけれども、ほんとうは次のページにデモのビデオを用意していたんですが、今ちょっとここでは見せられませんが、それぞれこの作業イメージの記事ごとにどういう記事なのか。タイトルなり、それから、それに対してどういうツイートが、この記事に対してどんなコメントがついているのかと、そういうコメントをネットから集めてきて、その情報を使って機械がランキングしている。これはずっと下のほうまでありまして、先ほどのチェックすべき可能性が高いものから、上

から並んでいる。それで、実際にそれをチェックするべきかどうかということを、人間がクリックしたりしながらチェックしていくということでございます。

実際に2018年9月時点のパフォーマンスというところで、ちょっとわかりにくいんですけども、スコア順にトップ、上から大体40%ぐらいまで見たら、ほんとうに疑義が呈されている言説の80%程度が見つかる。この時点では効率的な効果がそれほど高くはなかったんですけども、それでもとにかく現場に投入するということを今現在、行っています。

その例を次のスライド、その次のスライドですね。海外の研究動向はちょっと飛ばしまして、次の沖縄県知事選での活動について少しご報告させていただきます。これはF I Jさんが主体になって、昨年の知事選で実際にファクトチェックの活動をプロジェクト的にされたものですが、参加者のところに行きまして、6団体、6つのファクトチェック団体と、それから、F I Jが集めたサポートのメンバー26人から成るプロジェクトをやりまして、次のファクトチェックのプロセスですが、基本的には先ほどの仕組みを使いまして、F I Jの中で技術のサポートも使いまして疑義言説を集めてきます。疑義言説を集めてきたら、この6つのファクトチェック団体にそれを共有しまして、そのファクトチェック団体は、そのうち自分でほんとうにチェックしたいものを選んで実際に取材をして、記事化をしていくということをやるということでございます。

次のページの疑義言説を捕捉するということで、先ほどの仕組みを使いまして、実際に期間中に69件ほどの疑義言説を見つけることができました。次のページですが、それ以外に外部の参加者からの情報提供もありまして、全部で94件の疑義言説をF I Jとしてはファクトチェック団体に提供しました。

提供したのは、次のページの会員サロンというものがありますけれども、ネット上にその情報を配信しまして、ファクトチェック団体で共有する。そうすると、次のページ、先ほどと同じですが、ファクトチェック団体は、その中から自分たちでファクトチェックすべきものを選んで、取材をして記事化して、そしてF I Jからもファクトチェックのガイドラインを提供していますので、これにのっとってやられた記事については、F I Jのホームページでも公開するというような営みを、実際に技術と人が連携しながらやったというようにございます。

その次のページからは、実際に捕捉した疑義言説で、こうしたツイートの端緒情報から、次の琉球新報のようなファクトチェック記事が生まれてきました。ほんとうに両陣営さまざまな誤情報が飛び交っていたということがわかるんですけども、円グラフのほうまでいき

ますと、F I Jが捕捉した疑義言説というのはこれぐらいの、こういったもので、SNS上での疑義言説というか、言説にもかなり疑わしいものがあるということが、先ほどの仕組みで捕捉できていると。その発信者はさまざまでありまして、ただ、重要なことは、疑義言説だとF I Jとしては提供したわけですが、実際にファクトチェックをしてみると、そのうちの半数はそれほどというか、実際には正確あるいはほとんど正確であるというようなものもありますので、実際にファクトチェックをして確認するまでは、ほんとうのことは何もわからないということでございます。

そういうふうに先ほどの仕組みを使いまして、非常に広範なネットの情報の海から疑義言説である可能性の高いものを拾ってきて、その先は人間が全部やると。重要なことは、機械で全てフェイクかどうかということを判断するようなことは、現在の技術ではかなり難しいというか、ほとんど不可能だと思っております。ただ、一方、技術ができることもありますので、技術ができる部分を技術にやらせて、人間がやる部分とうまくデザインをしていくことが効率的に、それから、毎日これが続けるということができる重要なポイントではないかと考えています。

現在、この仕組みはF I Jの中で毎日動いていまして、毎日たくさんの誤情報、疑義言説を捕捉して、そして、さっきのクレームモニターという共有する仕組みの中でF I Jの協力団体、メンバーのようなメディアにそれを共有する。あるいは先ほどの技術の仕組みそのもののインターフェースをファクトチェック団体に毎日お見せする、使っていただくということまでできておりまして、その結果、今の最後の記事の絞り込みの絵の一番右側のところで、F I Jの中で日々○と×をつけているわけです。その○×を日々つけたこの情報自身がさらに機械学習のお手本になるわけで、それを日々機械学習をやっていくことによって、AIの技術自身もこの問題に対してよりセンシティブなモデルになっていくということで、今かなり精度が上がってきているというような状況でございます。

ということで最後、現状と計画ですけれども、F I Jが今この仕組みを継続的に使っていきまして、補助金等も利用してこれを継続していくような計画を立てております。F I Jのメディアパートナーにも無償提供をしております。という今、申し上げたようなことをやっているんですけれども、技術でまだできることはいろいろあると思いますので、そうしたことを今後も進めていきたいと思っております。

世界的な流れについてお話しませんでしたけれども、実際のファクトチェックの現場でこうして技術が動いていて、仕組みとして回っているというのは、世界的にもかなり先進的な取組ではないかと考えています。

以上です。ありがとうございました。

【宍戸座長】 乾先生、ありがとうございました。

それでは、次に資料3「フェイクニュース拡散の仕組み 計算社会科学の見地から」につきまして、笹原先生からご報告をお願いいたします。

【笹原講師】 名古屋大学の笹原と申します。私は科学の見地から、特に計算社会科学と言われる新しい学際的な科学の見地から、これまでフェイクニュースの拡散に関してわかっていることについて、幾つか知見を述べさせていただきます。

資料の3ページ目に自己紹介が載っていますがけれども、私がフェイクニュース問題にかわり始めたのは、実は2016年のアメリカ大統領選があった、まさにフェイクニュースの嵐が吹き荒れていた年です。アメリカのインディアナ大学というところで在外研究を行っていました。そこでエコーチェンバーという現象、この後出てくるんですけども、その数理モデルを始めたまさにその年に、フェイクニュースの問題が非常に顕在化しました。

次のページ、計算社会科学というところです。どういう学問かといいますと、新しい社会科学です。これまで人間の行動や社会現象というのは、なかなか定量的に研究することが難しかった。しかし、これだけデジタルテクノロジーが発達しまして、人々の一挙手一投足というか、行動がデジタル痕跡に残るような時代になり、そういったビッグデータを活用して、現実世界のプロキシとしてのデジタルデータを活用して、人間行動を定量的に分析することができるようになりました。あるいはウェブそのものを観察したり、実験や調査の場として活用して、非常に大規模に人間行動や社会現象を観察することができるようになりました。それから、そういったものを扱う、つまり複雑系を扱う数理や情報技術というものが発達してきて、これまで以上に定量的に人間行動や社会現象が研究できるようになりました。そういう見地から私は研究をしています。

6ページ目に行きまして、フェイクニュースの問題に関してです。これまでも何回も取り上げられたと思いますので、繰り返しません、非常に不確かな情報、あるいはあからさまにうそな情報というものが、特にネット、SNSを介して大規模に拡散をして、選挙に影響を与えたかどうかに関しては否定的だというお話が先ほどありましたけれども、少なくとも社会を、それから、政治を混乱させたということは間違いない。

フェイクニュースはどういう観点から科学するのがいいのか、というのはいろいろ議論はありますが、私が注目しているのはその次のページに書かれている拡散です。何が問題かというときに、あるいはどこを止めたらいいかというときに、やはり一番のキーになるのが拡散だと思います。SNSというのは非常に強力な情報拡散装置になっていて、誰でも

今は情報を発信できる。場合によっては、意図的なものもありますけれども、ひとたび人々のアテンションにとまると、それがシェアやリツイート、そういった共有機能を通してどんどん拡散されていって、こういったカスケードを起こすわけです。

その次のページに行きます。実際にこういう拡散を研究して、MITのグループですけれども、『Science』に論文が掲載されました。これはTwitter創業の年から2017年までの全てのツイートを、英語だけですけれども、対象として、ファクトチェック団体が、これはTrue、これはFalseということがちゃんとラベルづけてきたものに関してのみ、どういうふうに拡散したのかというのを調べたものです。パネルAのように、カスケードの深さに関しては赤いライン、虚偽の情報のほうが正しい情報よりも深く拡散する。それから、カスケードのサイズ、トータルでどのぐらいの人に届いたかということですが、横軸はログですので、桁違いにやはり偽情報のほうが拡散する。そして、その早さは、パネルEですけれども、同じ10という深さに達するまでにフェイクニュースのほうがはるかに早く拡散する。そんなことがわかりました。

それから、拡散されやすい話題というのもありまして、ここに幾つかトピックを挙げました。圧倒的に政治に関するフェイクニュースというのが、アメリカの場合は多かったということが報告されています。

次のページですけれども、誰が拡散しているかというのもアメリカの場合はだんだんわかってきています。フェイクニュースの共有に関しては、赤、黒、オレンジのライン、これが怪しい情報源からの情報をどのぐらい共有したかというのを累積で書いています。偽ニュースの共有の80%は0.1%のユーザによるものであると。あるいは高齢者や極右のユーザというのは偽ニュースを共有しやすい。そんなこともデータからわかってきております。

次のページに移って、もう一つ、拡散で問題なのはソーシャルボットと言われる存在です。これは2017年のカタルーニャ地方独立の住民選挙に関するものですけれども、点の一個一個がTwitterのユーザです。青の点々は人、赤はボットと言われる存在で、そのやりとりが線で書かれています。それはメンションだったりリツイートだったりします。ここから見るとれることは、赤のボットたちから非常に大きい丸、つまりハブユーザ、インフルエンサーに向かって情報が飛んでいるということ。つまり、このボットはハブを狙って情報を拡散している。そういう戦略があるということがわかります。

実はこのボットを凍結するというのは有効なアイデアだよ、ということを示した研究もありまして、それが右のパネルです。真ん中の緑の線というのが、ボットを取り除いたらどのぐらい偽ニュースは減るのかというのを示したものです。一番偽ニュースが減るのは、もち

ろんハブユーザを取ったときなのですが、ハブユーザはオーガニックなユーザであることが多いのですが、それを取ってはいけません。ボットを取るとセカンドベストで偽情報を減らすことができますよということを、この研究は示しています。

次のところに移ります。では、なぜ私たちはいとも簡単にフェイクニュースにだまされてしまうのかという認知的な要因ですけれども、1つは認知バイアスの問題があります。ウィキペディアに載っているだけでも180以上の認知バイアス、我々の認知の癖というものが知られています。簡単に言いますと、認知バイアスというのはある種の認知的なショートカットで、そういうものを使うことで自分に関係のある情報に素早くアクセスすることができます。それは進化的に備わった私たちの機能ではありますが、それには逆にフェイクニュースを信じ込みやすいという悪影響もある。例えば、直感や経験や思い込みなどです。

その認知バイアスの一種と数えられることもある、バックファイア効果というものが知られています。例えば右のパネルの一番下のラインを見ると、これはアメリカですけれども、自分の政治的イデオロギーとは逆の政治家のアカウントを1カ月フォローしてもらって、その後どういうふうにその人の政治観が変わったのかを調べました。しかし、基本的にやっぱり自分の世界観に合わない情報は見ないようにしますし、より自分の価値観に固執するようになります。結局、保守の人はより保守になる。そんなことが実験的にもわかってきています。

次に、もう一つ大きな私たちの認知の落とし穴は、社会的影響です。何かの意思決定をするときに、私たちは虚心坦懐に物事を見て、その中のベストを選ぶということは全然していません。私たちは他者から大きな影響を受けます。特に自分と似た他者から、そして、他者の感情からということです。そういったものがバンドワゴン効果、同調圧力あるいはホモフィリーや社会的ネットワーク、こういった言葉で知られているわけです。

これの1つの例をお見せします。これはやはりアメリカの例、横軸は単語の数です。銃規制や同性婚、気候変動、こういった話題に関するツイートに道徳的な単語が1、2、3、4、5、6と増えた場合に、リツイートされる確率がどう上がるかが縦軸です。道徳的な意味の単語だけだと横ばいなんです、そこに道徳的かつ感情的な単語であると、それはリツイートされる確率が上がっていくと。つまり、感情的なものに私たちは影響されやすいということを示しているわけです。

次に、今度は構造的な要因です。こういう認知的な脆弱性とプラットフォームの相互作用による影響で、1つはエコーチェンバーといわれる現象です。先ほど乾さんも発表されましたが、簡単に言いますと、似たもの同士だけでつながった閉じたこういうポンチ絵で書

いたような情報環境です。見たいものだけを見るとか、あるいはつながりたい人とだけつながると、こういったことが起こった結果、閉じた情報環境ができてしまう。一旦こういう構造ができてしまいますと、要するに自分の周りには自分と似た人ばかりですので、自分が多数派だと思いつまむようになりますし、自分と異なる意見は受けつけなくなる。そして、フェイクニュースの温床になるということになります。

本当にこういうことが起こっているのかというのを示したのが次のスライドです。これは我々のグループが、2016年のアメリカ大統領選中のツイートを分析したもので、青がリベラル系のユーザ、赤が保守系のユーザで、線がリツイートをあらわしています。このようにリベラル系はリベラル系とだけ専らやりとりをし、保守系は保守系とだけ。そして、異なるイデオロギーをまたぐようなコミュニケーションというのは、中に比べると少ない。こういうエコーチェンバーの様相を呈している。

こういうものを緩和しようと思ったならば、メカニズムを知らないといけない。次のグラフですけれども、動画を動かしていただけますか。簡単な計算モデルをつくってその原因を調べようと思いました。これは丸の一個一個がソーシャルメディアユーザで、最初の初期状態はリベラルから保守系まで、いろいろな意見の人がいるという状況です。それで、いろんなふうにつながっていると。自分と似た人から来た情報からは影響を受けて、自分の意見をちょっと変える。あるいは自分とは似ていない人はアンフォローしちゃって、誰か別の人をフォローする、ということを繰り返していくと、何が起こるかという、青い意見の人たちは、より青い意見になる。赤い意見の人たちは、より赤い意見になり、この青のグループ、赤のグループ、それぞれのネットワークのつながり自体もだんだん疎になっていって、最終的にはこれは分断されてしまう。そういうことが数学的には言えます。我々が現実で起こっている相互作用というのは、こんなに単純なものではないですが、そのエッセンスだけを取り出してきてモデル化すると、そもそもそういう分断されやすいという脆弱性というのが、このシステムの中に埋め込まれているという、そういうことを示すシミュレーションです。

次のスライドです。今週の「データの世紀」という日本経済新聞の記事で取り上げてもらったものです。最初にアメリカの民主党候補を300人フォローして、そういう状態でタイムラインの情報の多様性がどう上がるかを調べたものです。その中に、もし仮に10%異なる人をまぜた場合にどうなるかを示したのですが、自分とは異なる異質な人を多少でも入れたほうが、情報多様性が上がりますよということを示したものです。これは1つエコーチェンバーを緩和する方法なのではないかなと思っています。

2つ目の構造的理由は、フィルターバブルといわれる問題です。これはこの図を見ればお

わかりのとおり、無限に世の中には情報、特にネット上には情報があるのですが、私たちはそれを全部見るわけにはいかないので、アルゴリズムはそれをフィルタリングすると。その結果が、このアルゴリズムにフィードバックされますので、ここはポジティブフィードバックになっていて、どんどん自分の好みに合致したものが透過する、そういう世界ができて上がっていく。しかもそれが個々に起こるので、どんどん人々が見ている世界が分断されていくという現象です。

どのぐらいアルゴリズムで個人ユーザの属性を当てられるのかというのを示したのが次の図です。これはスタンフォード大学のKoshinskiらの研究ですが、例えばFacebookの「いいね！」のデータをたくさん集めてくる。どの記事に「いいね！」したかというのがわかると、いろいろなことがわかってしまう。例えば白人か、アフリカ系米国人か、あるいは性別、ゲイかどうか、レズビアンかどうか、そういった非常にセンシティブな情報まで、二択の問題に落としているのがみそなんですけども、非常に高い確率で当てることができる。

同じようなことを私たちのグループもTwitterでやってみたのですが、当てられないものもあるんですが、当てられるものもある。例えば、男か女か、デジタルネイティブかどうか、既婚かどうか、専門書を読むか、子供がいるか。あるいは、TwitterのデータでFacebook上に150人以上友達がいるか、つまり、ダンバー数を当てられるかという、6割強くらいの確率で当たる。そんなことがわかっています。

最後にまとめと、ちょっと私見を述べさせていただきます。まずはまとめです。フェイクニュース現象というのは、コンテンツの問題というよりは、むしろ情報生態系丸ごとの問題と捉えたほうがよくて、そういう視点を持つ必要がある。まず虚偽情報に対する要素レベルの脆弱性は、我々の認知レベルであると。それは認知バイアスだったり、社会的影響というものです。それから、そういった要素レベルの脆弱性がプラットフォームと相互作用することで、エコーチェンバーやフィルターバブルといった現象がマクロに起こって悪化する。そして、そもそも虚偽情報にだまされやすい背景になる情報過多の問題だったりとか、フェイクの自動化や巧妙化、ディープフェイクやボットの問題があって、ここはますます悪化している。これが現在の情報生態系という観点からまとめた、私なりのフェイクニュース現象の整理の仕方です。

次に、ポスト真実時代を乗り越えるためにどうしたらいいかというと、基本的にはメディアリテラシーを個々に高めるとするのはやはり必要なことですし、F I Jさんがやられているように、社会基盤としてファクトチェックという行為を浸透させていくというのはもちろん重要なことです。

全てのフェイクニュースを相手にする必要があるかという、そんなことはないと思っています。大体新しいフェイクニュースというのは、過去に起こったものを模倣してできるので、多分、典型的なパターンというはあると思うんですね。ですから、そういうものに関しては、やはりAIのテクニックやテクノロジーというのが効果を発揮するので、そういうものを使うというのは効果的だと思います。

2つ私見を述べさせていただきます。1つ目は、こういったフェイクニュースの科学からどんなところにコントリビュートできるかという、1つはメディアリテラシーですね。フェイクニュース現象はもちろん怖いんですが、なぜ怖いという科学的な根拠を知って怖がってほしい。ですから、フェイクニュース現象の科学的な要因、フェイクニュースの特徴や人間の特性、情報環境の特性、こういったものを例えば情報の教科書に入れるなり、こういったものを学ぶ機会というのがあってしかるべきかなと思います。

それから、先ほどのアメリカの研究では、高齢者ほどフェイクニュースをシェアしやすいという研究もありました。やはりデジタルネイティブよりはむしろお年寄りというか、高齢者の方に向けてこういったメディアリテラシーというものが必要になってくるのかなと思います。

2つ目は、ファクトチェックのプラットフォームで、こういうものがあつたらいいなと思っています。ファクトチェック・イニシアティブさんがやっているのは1つのやり方ですけども、例えばMe t a f a c tというプロジェクトがあります。これはファクトチェックのプラットフォームになっていて、こういうものを調べてくださいというのと、調べますというのがマッチングされていきます。裏でこういうファクトチェックをしているのは科学者で、いろいろな論文の情報に基づいて、特に健康と科学に関する情報に関してファクトチェックをしているんですが、そういった知見を貯める場所があります。日本でもやはりファクトチェックがこれだけ盛んになりつつあるので、少なくともファクトチェックされたものを貯め、再利用できる形でそれを誰でも使えるようにするというのは重要なこと。

このMe t a f a c tはいいんですが、問題は、やっぱり調べてくださいというものだけがたくさん貯まっていて、実はファクトチェックが追いついていないという現状があります。やはり寄附なりクラウドファンディングなり、そういったエコシステムとしてつながる必要があるのかなというふうに思います。

ちょっとオーバーしてしまいましたが、私の発表は以上です。

【宋戸座長】 ありがとうございました。

ご発表の最後になりますけれども、資料4「インターネットメディア協会の『信頼と創

造』への取り組み」につきまして、瀬尾代表理事からご発表をお願いいたします。

【瀬尾代表理事】 インターネットメディア協会代表理事の瀬尾と申します。よろしくお願い申し上げます。

私たちの協会で、信頼と創造についてどういう取組をしているのかということを、今日はご紹介させていただきます。

まずはインターネットメディア協会について、次めくってください。今年4月にできたばかりの団体です。会員はインターネットメディアあるいはそれに関連する事業を行う法人、団体または媒体を対象にしています。特徴については、ニュースメディア、エンタメメディア、コンテンツメディア、ニュースアプリなどのコンテンツメディアとか配信メディア、アプリケーション、さまざまなレイヤーの主要メディアが参加しているというのが特徴です。

次、めくってください。具体的な団体はこういう団体が加盟しています。加盟法人、媒体の一部です。BBCやBuzzFeedのようなニュース系のメディア、あるいはデジタル毎日さん、あるいはFNNのような既存メディアの中のデジタルを担っているメディア。そして、私が今SmartNewsという会社にいるんですが、SmartNewsあるいはGunosyのようなニュースの配信に携わっている会社、あるいはニュースだけじゃなくて、例えばカカクコムのようなコンテンツメディア、こういった形のさまざまなメディアが参加しております。この特徴は、これまでの業界団体のような、新聞協会や出版、雑誌協会の横ぐりの団体ではなくて、インターネットメディアをよくしようということを考えているさまざまなレイヤーの団体が集まって、平場で議論をして知見を共有することに意味があるというふうに考えたから、この団体をこういう枠組みで設立しております。

次をめくってください。インターネットメディア協会の目的ですが、第1に「信頼性」と「創造性」を通じた社会貢献をするということを目指しています。ネットメディアが何をもって社会に貢献できるのか。第1にはやはり情報が信頼できるということ。そして2つ目には、この多様なネット空間を通じて創造性を高めて、そのことによって社会に貢献していく。こういう集まりであることを目指しています。

2つ目には、先ほど申し上げたように、私たちは業界団体を目指しているわけではないです。業界の利益を代表する団体を目指しているわけではないです。社会とユーザの役に立つと、ここを目的にしています。そのために各レイヤーのメディアが集まって、自由闊達な議論と、そして、その議論をベースにした行動をしていこうということを目指しています。

3番目には、問題意識や知見を共有していこうと。ここでオープンな場として議論をして、ここで各知見を集めていこうということを考えています。このベースにあるのは、アメリカ

にOnline News Associationという団体がありまして、こちらはニュースにまつわる団体なんですけれども、例えばFacebookやTwitterのデジタルニュース担当者とか、あるいはニューヨーク・タイムズのデジタル担当者あるいは個人のジャーナリスト、デジタル専門のニュースメディア、こういった各レイヤーのばらばらのデジタルニュースにまつわる人たちが集まって議論をするという団体があります。そこで例えば成功事例、あるいはファクトチェックの取組とか、フェイクニュースの問題なんかを平場で議論していく。そのことによって各メディアの信頼性を高め、事業性も高めて社会に貢献していくという団体があります。

我々もこれはニュースだけに限る団体ではないんですが、インターネットメディア全体の社会貢献を目指すために、そういう各レイヤーで平場で議論できるということを重要にしていきたい、大切にしていきたいと思っています。

次めくってください。具体的な活動としては倫理綱領を策定します。これは我々のインターネットメディア協会は、どういうことを目指しているのか。ここに加盟するメディアは何を目指していくのかということ。そして、社会貢献と今、言いましたけれども、それは具体的に何かということを明示していこうと思っています。当然このメディアとしての信頼性あるいは人権の問題、あるいは言論の自由の問題、そういう基本的であるがゆえに大切なことを掲げていきたいと思っています。

それと同時に、発信者として信頼性向上にかかわる取組を勉強会で皆で共有していきたいと思っています。最新のそれぞれの取組であったり、さっき申し上げたようなアメリカの事例なんかも含めてここで共有して、各メディアの知見を高めていきたいと思っています。

さらには、この社会貢献のために各団体でいろいろなコンテンツのガイドラインなどに取り組むというところに対しては、我々がある種の知見を共有することでサポートしていきたいというふうに考えていきます。我々は団体としてガイドラインを掲げて、これを守るということを言うつもりはありません。要するに、我々としては各団体での取組をサポートしていきたいというふうに考えています。それと同時にJ I A A、広告系の団体ですね、こういったところとも情報交換をしていきたいと思っています。さらには、外部への働きかけとしては、読者リテラシーの教育など啓蒙活動へも取り組んでいきたいと考えています。

繰り返しになりますが、こういった取組のための交流の場にしていきたい、知見共有、交流の場をつくっていきたいと考えています。

次です。その中で信頼への取組ですが、特に信頼への取組ということでこういうことに取り組んでいます。倫理綱領、先ほど申し上げましたけども、これは今議論をしてほぼ固まります。近々発表できるのではないかと考えています。それを踏まえて、フェイクニュースや

ヘイト言論などにどういう対策をしていくのか、それぞれがどういう対策で取り組んでいるのかという知見を共有していきます。また、そういうことに対して各メディアの取組をバックアップしていきます。社会貢献にも積極的に行きます。

同時に、我々は個別の記事のファクトをチェックするようなネットの警察というところを目指しているわけではないです。各団体でご知見を共有して信頼するメディアの集まりを目指す、さらにはその先にブランド構築ができればいいなと思っています。

ネットメディアの信頼の向上というのを考えるときに、私たちはこういう図式ではないかと考えています。ユーザに対してメディアがあり、ネットの場合は広告収入も多いのでクライアントを想定しています。この三者それぞれが信頼向上のために取り組む必要があると思っています。その中でインターネットメディア業界としては、それぞれに何らかの働きかけ、影響を持っていきたいと思っています。我々自身、メディアについては自身で努力しなければいけない。そういう意味では倫理綱領をつくり、勉強会を開き、そして議論をして知見を共有する。そのことによって各メディアが信頼の向上のために取り組む必要があると思っています。

それと同時に、ユーザに対しても、例えばリテラシー教育の推進のような形で社会貢献をしていきたいと思っています。クライアントについては、これは先ほど申し上げた、例えば J I A A との連携や、あるいはクライアント側から我々インターネットメディア協会、J I M A という団体が信頼するメディアの集まりであるということを認知してもらえれば、こちらのブランド認知によってクライアントが我々のメディアを評価してくれれば、結果的に信頼するメディアの価値が向上して、悪貨が良貨を駆逐するわけではなくて、良貨が悪貨を駆逐するようなメカニズムができるのではないかというふうに考えています。

特に、我々 J I M A というブランドを認知していただきたいと思うのは、クライアントが個別の記事の判断をしていくというのは、やっぱり言論の自由や報道の自由から考えると、ある意味危険な部分もあるんです。そういう意味では、我々のような団体をブランド認知していただけるのがありがたいと、方向的には思っております。

次です。リテラシー教育、皆さんご存じかもしれませんが、どういうことをしているのかということを簡単に紹介させていただきます。

こちらの写真にあるのは、インターネットメディア協会の設立記念シンポジウムというのを先日開催しました。このときにリテラシー教育のセミナーも同時に開催しました。合わせておよそ 300 人の方が参加してリテラシーについて学んでいただきました。こういうセミナーだけではなく、ホームページ、メルマガを通じてさまざまな情報発信をして啓蒙活動も

していこうと思っております。

その一例なんですが、次をめくってください。具体的にこういうことをやっています。

リテラシー教育というとなんて難しいようなんですが、こういうことをやっているんです。これは小学校でやるプログラムの一つです。こういう事例を子供たちに見せます。「先生に 卒業式のことを聞こうとしたら 返事をするのがいやだったのか なかなか顔を合わせてくれなくて 私たちに逃げるように 体育館のうら口から職員室に コソコソ行ってしまった」。めちゃくちゃ嫌な先生という印象を受けますよね。じゃあ、子供たちがこういう話を聞いて、さあ、この中で事実は何でしょうか、あるいは事実でない印象や感想はどこでしょうかということ聞くわけです。

そうすると、「先生に卒業式のことを聞こうとした」のは事実です。「返事をするのがいやだったのか」、これはわかりません、臆測ですね。「なかなか顔を合わせてくれなくて」というのも、これも憶測です。「私たちに逃げるように」、これもどういう意味だったか意図はわからない。そういう意味で言うと、体育館の裏口から職員室に「行ってしまった」のは事実だけど、「コソコソ」は印象です。

というわけで、次をめくってください。事実としては、「先生に 卒業式のことを聞こうとしたら 体育館のうら口から職員室に 行ってしまった」というのが事実であると、こういうことを教えています。我々は、メディアリテラシーというのは単なるメディアの見方だけではなく、情報に対する感度、見方を子供たちがちゃんと自分たちの頭で考えられるようにしようという教育を提供していこうと思っています。

次はフェイクニュース対策への私見です。こういうことを私は思っています。フェイクニュース対策というのは、これは難しい問題があります。まず第一に、定義が曖昧であるということ。もともと、フェイクニュースというのはネットニュースのことではなくて、ランプがCNNのニューヨークタイムズ批判に使ったわけですね。「ラベル貼り」であったということ。あるいは、一般に意図による分類が可能だと、つまり、誤情報と偽情報に分けられるといいますけども、発信者の意図というのは、実は極めて難しいところがあります。意図だけで分類するのは、これはわからない。そういう意味では定義が曖昧になりかねないということです。

次です。対策はいろいろあるんですが、即効性はそれぞれないです。ファクトチェック団体にやってもいい試みだと思うんですが、これも全てが対応できるわけじゃない。ビジネスモデルの撲滅も、判定がそれぞれ難しい。プラットフォーム業者による規制も、その介入性、中立性も問われますし、数が多くて全てに対応できるかという問題もあります。行政による

規制は、当然、言論の自由や報道の自由への侵害になる危険があります。そういう意味で反対です。つまり、それぞれいろんな対応がありますけども、即効性はなく、これを強くしようとするとな劇薬になって大変な影響を持ちかねないという意味で、少しずつを組み合わせた対応しかないのではないかと思います。

その中で私が考える行政が取り組むフェイクニュース対策というのは、こういう取組をしていただければというふうに考えております。原則としては、やっぱり自主的な取組を中心に考えるということ、やっぱりマーケットのメカニズムの中でビジネスモデルとして信頼を尺度にして市場で淘汰をされていくというのが、やはり一番理想的な考えだと思います。ここをぜひ目指してほしいと思っています。そういう意味で、行政の規制については、言論の自由も含めて反対です。現行法の中で十分対応できる部分があるので、そこでぜひ考えていただきたいと思っています。

じゃあ、具体的に何をしたらいいのかということなんですが、まずは実態調査、これは大学や研究機関、民間の研究機関でも取り組んでいますけれども、データの収集も含めてかなりコストがかかって十分ではありません。フェイクニュースの実態はまだわからないことがある。影響もよくわからないことがあります。ここはぜひ行政でバックアップして民間、大学など、研究機関を支援して、実態をまず調べていただきたい。その上で対策を考えたいです。

さらに、情報公開、透明性の拡大を進めたい。これはファクトチェックするときにも、そのもとになるデータに当たるのが大変です。これは行政の情報だけではなくて民間も膨大なデータがあるんですが、それぞれデータの保有の形式とかスタイルが違って、とりまとめるのが難しいような状況です。これも例えば、行政主導で官民のデータの形式をそろえることによってファクトチェックがしやすくなる。普通の方が直接生データに当たれるような状況をつくるということが全体の信頼の向上につながると思います。

そして最後、情報リテラシー教育の導入です。これはぜひ教育現場で子供たちにリテラシー教育をしてほしいと思っています。海外でも導入事例があります。これも行政がプログラムをつくるというよりは、例えばこのインターネットメディア協会もそうですし、私が今いるSmartNewsもそうなんですが、リテラシー教育は実績があり、民間のプログラムがあります。海外でも民間のプログラムを中心に行政が指導してリテラシー教育に取り組んでいる事例が多々あります。ぜひ、この部分については推進をしていただければと思っています。

最後です。私たちはよくインターネットメディア協会で、インターネットメディアはどうなるのかと聞かれるんですが、これについてはこういう見解です。これは僕の個人的な見解

ですけれども、メディアの未来は率直に言ってどうなるかわからないと。スティーブ・ジョブズだって一回会社を首になったぐらいなので、これがわかったら苦勞しない。けれども、私たちとしてはどういうメディアの世界にしたいか、この世界観と意思は明確に持っておく必要があると考えています。

その柱になるのは、私たちのインターネットメディア協会の取組では信頼と創造性、これをもって社会に貢献するというネット空間を目指していきたいと思っています。

以上、私の発表はこれで終わらせていただきます。

【宍戸座長】 ありがとうございます。

それでは、ただいまの4つのご発表についてご質問、また、全体を通じて構成員の間での意見交換を行わせていただきたいと思いますと思いますが、ご質問等に関して、私からあらかじめ留意いただきたい点がございます。

1点目は、時間の関係上、なるべく多くの構成員の皆様からご発信、ご発言をいただきたいと思いますので、手短にご発言、あるいはご質問をいただきたいと思いますと考えております。

2点目に、ご質問の際にお越しいただいた有識者の方全員に、広域制圧型で全員にお伺いしますとか、そういうことは時間がかかりますので、誰に対して何をお伺いになっているのか、そこを明確にしてご発言をいただければと思います。

それでは、挙手をお願いしたいと思います、いかがでしょうか。

では、生貝構成員、お願いいたします。

【生貝構成員】 貴重なお話をいただきましてありがとうございます。

乾先生に1つと、成原先生に2つ、短くお伺いしたいと思います。

まず、乾先生の資料の最後の34ページの上から2ポツ目のところに、お時間がなくてご言及いただけなかったかもしれないんですが、Googleなどの大手プラットフォームとも連携を促進しという記述があることにに関して、やはりこの分野でこういった取組を進めていくに当たって、大手プラットフォームの間で今どのような協力が行われていて、そして、これからどのような協力が行われていくことが望ましいのかということについて、もしご意見があれば教えていただきたいというのが乾先生へのご質問です。

【宍戸座長】 では、一回切って、お願いできますか。まず乾先生、お願いします。

【乾教授】 ありがとうございます。時間の関係で申し上げられなかったんですが、非常に重要な点だと思います。

例えば、Google等も含めてプラットフォーム事業者がクレームレビューという枠組みを共有して、ファクトチェックをした情報等のネットでの共有の仕組みというのを、世界的な規

模で積み上がってきているものがございます。

例えば、ファクトチェックの情報がありましたら、特定の特殊なメタ情報のようなものをそこに埋め込んでおくと、例えばGoogleが認めた団体が発信しているファクトチェック情報のようなもの、そうしたものに特定のメタ情報を埋め込んでおくと、それがプラットフォーム、例えば検索の中で反映されるというような仕組みが国際的に流通してきているということがございます。

ですので、残念ながら日本の団体でまだそういう活動を実際にできているという事例がまだほとんどございませんので、日本のファクトチェック団体もそうした仕組みの上に乗っていきけるようにしていく、その体力、それから技術的な面も含めて、していく必要があると考えております。ファクトチェックイニシアチブのほうもそうした活動が進んでいくように今取り組んでいるところではございますけれども、さまざまな支援が必要だというふうに考えています。

ありがとうございます。

【生貝構成員】 ありがとうございます。

それでは、成原先生にも1点にいたします。全体として国家によるフェイクニュースへの規制にかかわる議論というものを教えていただきました中で、やはり他方でこれだけプラットフォーム方々もフェイクニュースの対策というもの世界中から求められると、いろいろなアカウント等を消すようになって、重要なアカウントが、全く良い情報を発信しているはずのアカウントが突然消えたりするというのが、日本でも少しずつ起こり始めているような気がいたします。

そのようなときに、例えば彼らに対してどういう基準で削除しているのか、あるいは凍結をしているのかということをお願いしていくことが、やはり自由な言論や、あるいは政治的な表現の重要性に鑑みて必要なのではないかという議論は、こういった議論はアメリカではあるのかということについて伺わせていただければ幸いです。

【成原准教授】 ありがとうございます。

ご質問に関連する問題といたしまして、報告の中でも少し申し上げましたとおり、アメリカの保守派の中には保守派の言論がプラットフォーム事業者により不当に削除されている、アカウントが削除されているという不満が一部にございます。

一方、プラットフォーム事業者の主張によりますと、あくまでも政治的な思想、考え方、見解に基づいて削除をしているのではなくて、ヘイトスピーチに当たるだとか、あるいは暴力を扇動しているとか、各々のプラットフォーム事業者の定める基準に従って削除をしてい

るということです。

また、このような保守派の問題意識を受けまして、私の報告でも紹介しましたように一部の議員からCDA230条の免責を受ける前提としてプラットフォームの政治的な中立性を求める法案も提出されております。この法案では、FTC（連邦取引委員会）がプラットフォームの公平性が担保されているか認証するという事になっているんですけれども、国家の機関に公平性の判断を委ねるということで、やはり表現の自由の観点から問題ではないかという指摘があります。

そこでやはり生貝先生がおっしゃられるとおり、各プラットフォーム事業者において、各々の基準に即してどのような理由によりアカウントを削除したのか、あるいはコンテンツを削除したのかということを自主的に公表していくというのが表現の自由の観点からも、あるいはそのプラットフォーム上の自由な情報流通を実現する観点からも望ましいのではないかと。公平性と透明性の両方にかかわってくる問題ですけども、公平性については何が公平なのかという難しい実体的な問題がかかわってきますので、まずは手続的に透明性を確保していくところから始めていくのか大事なのではないかと考えております。

【生貝構成員】 ありがとうございました。

【宋戸座長】 ほかにいかがでございましょうか。

それでは、まず森構成員、お願いします。

【森構成員】 ご説明ありがとうございました。私も成原先生に米国の憲法の議論のことお尋ねしたいと思います。今のその中立性みたいなお話ですけども、確かに情報自体の問題とかアカウントの問題ということもあると思うんですが、他方で、先ほど来ご説明のありましたフィルターバブルの問題、エコーチェンバーの問題もあるわけで、受け手のほうに、何と言うんでしょう、従来であればいろんなものを見て判断するということができただけですが、そういうことができなくなっている。それは自然発生的に起こってやむを得ない部分もあると思うんですけども、そうでなくて、例えばその政治的広告であっても、ターゲティング広告によって出されることがあると。悪いシナリオとしては、高齢者をターゲットとすとか、そういうことだってできるわけですので、そういった観点から、ある種、人工的、人為的に作り出されたフィルターバブルみたいなものが、それは全く自然なものではないわけで、それが、知る権利、受け手のほうの権利との関係で問題になるのではないかと、そういう議論があってしかるべきだと思うんですけども、そういう話というのはないものなのかお教えていただければ幸いです。

【成原准教授】 アメリカにおいても、先生がおっしゃるようにエコーチェンバーによっ

て一定の見解や思想だけに包まれたような情報環境が構築されることは、知る権利を実質的に充足する観点から問題だという議論はございます。ただ、判例あるいは通説のような立場にはなっていないのが現状です。インターネットとは少し文脈が異なるんですけども、アメリカでも放送については、1987年までは、放送局に、公共の関心事に関する論点について対立する見解を公平に扱うことを義務づけた公平原則が存在して、日本の放送法の番組編集準則におおよそ相当する内容をもっていたんですけども、これも廃止されてしまいました。その後、政治的分極化が激しくなったということも受けて、今日では公平原則を再評価する議論も生じております。ただ、放送のみならずインターネットにまで公平性を求めることが表現の自由の観点から適当かどうかについては、アメリカでも議論が分かれているところがございます。

日本について申しますと、先ほども申したように放送法4条で番組編集準則が定められていて、その中で、政治的な公平性などについて定められており、それを踏まえ放送局が自主的な取組をしております。先ほども申し上げたように、放送は、ソーシャルメディア、インターネットが発達したアメリカにおいてですら依然として大きな影響力を持っていて、これはおそらく日本でも同様の傾向を見てとることができるでしょう。そうだとすれば、エコーチェンバーを防ぐという観点からは社会全体の情報環境のあり方が重要になってきますので、仮にインターネットの情報流通に公平性を求めるべきではないとしても、放送が公平性を実現するプラットフォームとして重要な役割を果たしていく可能性というのはあるのではないかと考えております。

【森構成員】 ありがとうございます。確かにそれはまさにおっしゃるとおりだとは思いますが、ただ、フィルターバブルなりエコーチェンバーなりは、実際に我々の目前で起こっている問題ですし、我々はそのスマートフォンの画面を見ている時間が長くなっているわけですが、シンプルに言ってしまうと、例えば政治的広告についてのターゲティング広告をやめると、そろそろマス広告のみにするという事にすれば、フィルターバブルを使った誘導的なことというのは防げるんじゃないかと思うんですけど、もちろんそれは極端な意見だということはわかっていますけど、そういう議論はないんでしょうか。

【成原准教授】 もちろんございます。関連する論点が2つほどあると思っております、まず、公平性を求めるべきか否かという論点と、ターゲティング広告を規制すべきか否かという論点があるかと思います。仮にプラットフォーム事業者に公平性を義務づけなかったとしても、プライバシー、個人情報保護の観点から、プロファイリングやターゲティング広告を規制するという選択肢はあるわけでございまして、こうした議論はアメリカにおいてもプ

ライバシー研究者を中心にで活発になっております。さらに一歩進んで、個人のプライバシーだけではなく、健全な民主主義プロセスの維持・促進といった社会的・公共的価値も勘案して、プロファイリングやターゲティング広告の規制の在り方を検討していくことも考えられるはずで、アメリカでもそのような議論は盛んになっているところと認識しております。

【森構成員】 ありがとうございました。

【宋戸座長】 よろしいでしょうか。それでは宮内構成員、お願いします。

【宮内構成員】 乾先生に質問が1件と、瀬尾先生に1件ありますので、それぞれお聞きします。

まず乾先生、これは機械学習を利用しようということで、このチェックすべき記事は濃縮していくということになりますね。こういうのは、機械学習の中では非常に筋の良いアプリケーションだと私も思っていて、ぜひこれを進めていきたいと思っています。今のところ、まだ少し性能としては不十分なところもあるというお話もありましたけど、これは今後の見込みがどうかというのを教えてほしいんです。つまり、データが集まればかなりいけるのか、やはり、今までデータが集まってきてだんだん少しずつはよくなっていると思うんですけど、やはり相当限界があるのか、このあたり、今後どういうふうに進みそうなのかというご見解を伺いたいと思います。

【乾教授】 はい、ありがとうございます。性能がどこまで伸びるかの予測というのはなかなか難しい、問題の性質が、データの性質なり偏りがどうなのかというのは実際にやってみないとわからない。それから、これはほんとうに現実のデータを見ているので、研究室の中できれいなデータでやっているわけではないということで、さまざまに、今実際に、当初予想しなかったようないろんなデータが出てくるわけです。ですから、まさに生きているデータを扱っているという意味で、技術がどれだけできるのかということの予測がなかなか難しいというところがございます。

技術に頼り過ぎない設定というのがとにかく重要であるというふうに思っておりまして、今回の場合ですと最終的に確認するのは人間ですし、それからチェックするのも人間である。その人間がやる部分をもう少し効率よく、かつ網羅性があるように担保できるようにしていく、その程度問題として少しずつよくしていくということはあると思うんですけども、それがどこまで行くかというのを今の段階で予測するのは難しい。そういうことでございます。

【宮内構成員】 桁が一つずつ上がっていくごとにどのくらいよくなってきたかというのは、それなりの傾向は……。

【乾教授】 今まさに、この6カ月ぐらいのデータがまとまってきていますので、それを定量的に評価するということを今まさにやろうとしていまして、ちょっと今日は間に合わなかったんですけど。

【宮内構成員】 そういうものを進めておられて、これからということですか。

【乾教授】 そうですね。

【宮内構成員】 わかりました。どうもありがとうございます。

もう1点、瀬尾先生に教えていただきたいんですけども、瀬尾先生の13ページのところに、特にここに書いていることは非常に私も賛成なんですけど、情報リテラシーの教育に關しましてちょっと教えていただきたいことがございます。

先ほどの例にあったとおり、どこまでが事実なのか、どこまでが推測なのかということを区別すると、非常に重要だと思うんですが、私も論文とか書くときにすごく教育されたのが、事実と推測を分けて書けと。どちらかわかるように書けというのは、発信側にとってもすごく重要なことだと思っております。

実際には、マスコミのニュースにしても個人のツイートにしても、事実なのか推測なのかがよくわからないものってすごく多くて、実は誤情報が伝わってくることの温床になりやすいんじゃないかと思っているんですけど、このあたりについて、例えば教育の観点とか、そういうところはいかがお考えでしょうか。

【瀬尾代表理事】 とても重要なご指摘だと思います。リテラシー教育という場合に、一般の消費者教育とまた少し違うところは、要するに情報の受け手だけではなくて、今のソーシャルの時代は情報の発信者になり得るということなんです。やっぱりフェイクニュースの原因の一つは、例えばビジネスモデルの問題とか、あるいは政治的意図によるある種の工作的なものもありますけども、やっぱり拡散する大きな原動力になっているのは、一般の方のただの好奇心とか、悪意なき拡散というのが一番大きな問題になっています。

そういう意味で言うと、先ほど申し上げた意図だけで見分けるのは難しい。要するにフェイクなんだけども、善意で拡散しているケースもあるというケースがありますよね。そういう意味で言うと、リテラシー教育が重要なのは、受け手だけではなくて自分が発信するときはどうなのか、やはりその事実と推測、あるいは感想、感情を分けて書かなきゃいけない。そこを区別して発信しなければいけない。そこを教育していくことが極めて重要だと思っております。そして、その先にあるのはただのメディア教育だけでなく、やっぱり情報の受けとめ方、発信の仕方の教育につながっていくと思うんです。

実はこの問題は、僕はぜひ教育でやっていただきたいと思っているのは、単なるメディア

リテラシーでなくて、例えばほかの金融リテラシーとかとも通ずる、今日本に一番求められている答えのない世界の中で自分の頭で考えるということ求める教育に行き着くのだと思っています、その入り口として、受信者、同時に発信者としてのリテラシー教育というのは極めて重要だと思っています。

【宮内構成員】 どうもありがとうございました。

【宋戸座長】 ほかにいかがでしょうか。寺田構成員、お願いいたします。

【寺田構成員】 瀬尾さんにお聞きしたいと思います。ネットでメディア間の情報共有が進むというのは、やっと進み始めたのかということ非常に心強い限りなんですが、その一方で、こういった取組というのは一歩間違えるとメディアの多様性といったものを奪いかねないというところもありますので、そういったところの対策とか、そういったことをどう考えていらっしゃるのかということと、先ほどからお話が出ていますが、結局どのようにそのデータが拡散していったりとかというときに、発信者であるメディアサイドがこういった技術的な対応、トレーサビリティの対応であったりとか、極端な話、データにタグを埋め込むとか、そういったことも含めたフレームワークのようなものというの、今後、検討の課題とかの中にあるんでしょうかという2点になります。よろしくをお願いいたします。

【宋戸座長】 では、瀬尾代表理事、お願いいたします。

【瀬尾代表理事】 まず最初の問題なんですが、実はこの団体を立ち上げる時も、そこが一番議論されたところです。私たちは、まず一つには、この団体自体が掲げている信頼性と創造性という問題ですよね。で、この創造性を支えているのは多様性なんですよね。多様性を支えるためには、各媒体の独自性、あるいは発信力というのが極めて重要だというふうに考えています。

その立ち上げのときにも、実はこの団体を立ち上げるときに、最初に立ち上げまで2年ぐらいかかっているんです。その中で、例えばガイドラインのようなものをどういうふうにつくっていくのか、それをどういうふうに取り組んでいくのかというのをかなり時間をかけて議論しました。我々がたどりついた結論、今年の4月に立ち上げるときには、要するに表現の自由や言論にかかわることを、各媒体で中心になって独自で取り組んでいきたいと思います。先ほど申し上げたように、全体的に我々が何をすべきか、信頼性を向上させるべきとか、あるいは基本的人権や差別についてどう取り組むかということは倫理綱領で定めていきたい。それに基づく各媒体の表現や言論については、やっぱり先ほど言った警察のような取り締まりではなく各媒体で取り組んでほしいというふうに考えています。そこを大事にたてつけた団体になっています。とはいえ、各媒体でそういう取組は進めてほしいと思っています。

ので、勉強会の中で実際どういうことができるのか、そういう議論をしていきたいと思っています。

後段のご質問なんですけども、技術的な対応とか、そういう部分に関しては、私たちとしてはこの勉強会の中で議論していきたいと思っています。この団体に意味があると思っているのは、いわゆるこれまでの業界団体的な団体ですと、横的なつながりで、まさに例えば技術的な問題とかデータの的な問題、あるいはテクノロジーの問題とか、媒体だけでは対応できない、従来型の媒体では対応できないというケースもあったわけです。なかなかそういう事例も共有できないというケースがあったので、我々の場合、例えばニュースアプリの会社のような、かなりテクノロジー技術を持った会社もあります。あるいはデータに接している会社もあります。そういう会社がこういうコンテンツメディア企業と一緒に議論をすることによって新しい可能性を見つけていきたいと思っています。そういう議論の場所、あるいは勉強会をこれからつくっていききたいというふうに考えております。

【寺田構成員】 ありがとうございます。

【宍戸座長】 ありがとうございました。ほかにいかがでしょうか。

大谷構成員、お願いします。

【大谷構成員】 ありがとうございます。笹原先生の資料3のところで、最後のほうにご説明をいただいている項目がございまして、特に36ページのところ、アルゴリズムは悪さをしないのかといったところについて、2015年の研究の結果だということなんです、これが2019年の現在においても変わらないのか、あるいはこのグラフの見方などについてご教示いただければと思っています。

実は、これまでプラットフォームの方々のヒアリングなどをさせていただく中で、プラットフォーム業者がとても大きいと思っているんですが、どこを期待していいのか、あるいはここは期待できない部分なのかという整理がやっぱり必要だと思っております、その参考になる情報だと思っておりますので、よろしくお願いいたします。

【笹原講師】 ご質問ありがとうございます。これは補足資料としてつけたものなんです、先ほどからエコーチェンバーやフィルターバブルというのはやはり問題だということは言われている、問題だというポテンシャルはあるんですが、実際はどうなのというのをはかされたことは実はなくて、フィルターバブルに関しては、おそらくこの2015年の研究が、アルゴリズムはそこまで悪さをしていないのではないかとすることを初めて示唆する結果を出した論文になります。このグラフの見方ですけども、これはFacebookの研究者が、ラダ・アダムックという女性のデータサイエンティストですけども、そのグループが中心となっ

て、Facebookのデータを使って、アメリカの場合はユーザが、自分がリベラル側なのか保守側なのかということ的主張する項目があるので、その情報を使って、もし仮に自分が保守側だった場合、これ、何もしなくても40%強の確率で自分のイデオロギーとは逆のイデオロギーの情報が、40%ぐらいの確率で見える確率がある。それはリベラルの場合も同じで、40%強の確率で自分のイデオロギーとは逆の情報を見る可能性があるということを示しているんですが、次に、友人がシェアした情報は優先的に見えるようになるので、そういう過程を経ると、一気にその反対側の情報が入ってこなくなって、それがぐっと落ちていくのがこの部分です。

その後、どの情報を上に見せるかというアルゴリズムの操作が入るんですが、その後は、保守派もリベラル派も、その反対の情報が見られるという確率が実はほんの少ししか下がってなくて、この差分だけを比較してみると、実は自分の意見とは逆の意見を見えなくしているのは、あなたが誰とつながっているかというソーシャルネットワークのエフェクトなんですということを示したわけです。

これは、Facebookの研究者がFacebookのデータでやっているの、私たちは再検証のしようがないんですが、もちろんオーネストにやっているとすれば、現段階では、実はFacebookのアルゴリズムはそこまで悪さをしない可能性はあって、今どうかというのは、まだ実際にデータがないと何とも言えないんですが、その後、Facebookはいろいろ改善をしますので、大きくこの状況は変わってないのではないかと思います。ですから、フィルターバブルの影響というのは、実はポテンシャルとしてはあるんですが、今のところはそこまで深刻でないのかなという感じを受けます。

【大谷構成員】 ありがとうございました。

【宋戸座長】 ありがとうございます。ほかにいかがでございましょうか。

では、森構成員、お願いします。

【森構成員】 ありがとうございます。今の話で重ねて、笹原先生に教えていただければと思います。

ニュースに関することとしてはそうだと思うんですけど、広告は違うじゃないかという印象を受けるのですが、それはそういう理解でよろしいでしょうか。

【笹原講師】 これは広告とかそういうのではなくて、オーガニックなユーザが自分の身の周りについて発信した情報に関してのみトレースしていますので、ターゲティング広告みたいなものはここには含まれなくて、それはもちろん分けて考えないといけないと思います。

【森構成員】 ありがとうございます。

【宍戸座長】 ありがとうございます。ほかにいかがでございましょうか。

では、生貝構成員。

【生貝構成員】 今のお話に関連して、笹原先生に、もしご意見があれば教えていただきたいのですが、先ほど、まさに最後の36ページのところで、Facebookの研究者がFacebookのデータを使ってこうした検証をして、そのデータ自体は今は手に入らない状況にあるといいましたときに、やはりこの問題というのはすごくEvidence-basedで、まさに成原先生がおっしゃったとおり、しっかりと事実を検証しながら取り組んでいかなければならない。そのときに、やはりデータがなければ検証ができないし、データがあれば検証ができる。どうもそのデータは、いわゆるプラットフォーム事業者と呼ばれるところにたくさんあるらしいということが何となくわかってきたのですけれども、もし例えば、こんなデータが手に入ればもっと豊かなEvidence-basedの検証ができるのにといったデータがもしあれば教えてください。

【笹原講師】 やはりプライバシーがかかわる問題ですので、匿名化の操作を施しても、実は特定できてしまったりという問題があるのでなかなか難しいですが、少なくとも、そういうことがもっと安全にできるようになったら、できるだけ丸ごと欲しいというのが正直なところです。少なくとも、こうこういう結果を再現できるぐらいの量のデータは欲しいなと思います。この具体的なグラフに関しては、もちろん加工されたデータは公開されています。ですから、リプロットはできます。ただ、この操作を行うに当たっての前処理はどんなふうにしたのかというのは、やっぱり前処理の仕方が変わってきますので、そのあたりのデータも欲しいかなと思います。

【生貝構成員】 ありがとうございます。

【宍戸座長】 ありがとうございます。さらにいかがでございましょうか。ご質問、あるいは今までの質疑応答も踏まえたご意見でもよろしいのですけれども。

もしなければ、私から瀬尾代表理事にお伺いしたいことが1点ございます。

それは、いわゆる見出しの問題になります。一般にこれまで、雑誌あるいは新聞の記事といった場合に、見出しというものが非常に重要である。それは記事の中身を、見出しをもって理解するし、あるいは見出しの書き方によって、その記事の印象が大分変わってくるということもございますし、また、中吊り広告などがそうですけれども、雑誌あるいは新聞の購買といった行動に読者を結びつけるという意味でももともと重要だったものですが、現在、インターネット上で、いわば見出しだけがニュースサイトなどで流れていって、それでイメージみたいなものがつくられていたり、その下を読まない、あるいは十分に読まな

いで何かイメージがつくられて、それこそフェイクニュースであるというような批判が起きたりすることがある。こうした問題が幾つかあると思うんですが、インターネットメディア協会、または瀬尾さんのこれまでのジャーナリストとしてのご知見も含めて、何か改善点とか、発信者側もそうですし、プラットフォーム側、あるいは利用者側、どこでもいいのですが、お気づきの点があれば教えていただきたいのですが、いかがでしょうか。

【瀬尾代表理事】 いわゆるネットでも釣り見出しというのがよく話題になりますよね。ご指摘のように、見出しが人を引きつけるというのは便利な面もあるし、問題の面もあります。そもそもネットだけの問題かというところでもなくて、昔から週刊誌でもそうですし、テレビでも続きはCMの後で的な、惹起というのはよくありますよね。ネットの場合に問題なのは、それが極端ではないかと。あるいは全く中身が違うというケースもある。だから、これまで持っていたマスメディア的なものの悪い面を拡大しちゃっている面があるんじゃないかなという認識は広がっていると思います。これは特に定量的なデータは僕も持っていないんですけど、感覚的にはそうなのかなというふうに認識しています。

この問題についても、僕ら、インターネットメディア協会の中で議論して、当然、問題認識があればそれをどういうふうに改善していくのか、あるいはそれをどういうふうに取り組んでいくのかということは、語るテーマの一つになると思います。それは先ほど申し上げましたように、僕らは業界団体として利権を、利益を代表するわけじゃなくてユーザを守る、社会に貢献するということを目指しているので、その観点から取り組んでいきたいテーマだと思っています。

なぜ、ネットの場合、それが既存のメディアよりも拡大する可能性があるのかというと、いわゆるこれはビジネスモデルの問題でもあると思っています。そのネットの中の一部広告には、要するに記事の中身に関係なく、いわゆるPV、クリックされたことによって稼げるというモデルがありますよね。これがフェイクニュースの温床にもなっている部分もあります。そういう意味で言うと、ビジネスモデルの中で改善していく、要するに信頼というのを尺度に記事の価値がついて広告の値段もつくようになれば、当然、信頼できるメディアに広告を出すことが企業のブランドにとってもブランドアップにつながるわけですし、あるいは企業の社会的貢献としても社会を守ることにつながるわけなので、そういう尺度になっていくことによって、自然と市場メカニズムの中でクリアになっていくのが一番良いのじゃないかと思っています。

特に釣り見出しの問題は、僕はメディアにとっても重要な問題だと思っているのは、仮にそういうことをしてお客さん、読者の方が来たとしても、要するに、見出しと中身が違

いますという、つまり、中身がうそかどうかは別で、見出しと中身が違いますという場合、来た方はメディアに対する信頼性を失うわけですね。そういう意味で言うと、メディアの信頼を毀損していくビジネスモデルであるわけで、タコが自分の足を食っているようなビジネスモデルになるので、釣り見出しの問題は、やっぱりメディアの信頼という観点からも取り組んでいくテーマであるというふうに考えております。

【宋戸座長】 丁寧ありがとうございました。

さらにご質問、ご意見、いかがでございましょうか。

では森構成員、お願いします。

【森構成員】 ありがとうございます。私も瀬尾さんに今の釣り見出しの関係でお尋ねしたいんですけど、釣り見出しといいますか、それネットメディアの場合、なかなかマネタイズということでいろんなご苦労があるろうかと思うんですけども、フェイクニュースとはちょっと違うのかもしれないんですが、ちょうちん記事みたいなものといいますか、すいません、言い方はよくないのかもしれませんが、何ていっていいのかわからないのですが、例えば、企業サービスなり商品なりを紹介して、それはただただ何がしかのブレイクスルーであって社会に伝えるべきことであるということもあるかもしれませんが、場合によってはスポンサーであったり、さらに言いますと、より複雑な形としては、将来広告を出してくれるかもしれないねみたいなこともない話ではないと思うんですけども、そういうことについては、その何かご議論みたいなことはありますでしょうか。

【瀬尾代表理事】 ネットメディアの中で、いわゆるその広告と記事の区別がつかないような悪質な記事、あるいは広告があるという問題はかつて指摘されてきて、これも大きな問題の1つだと思っています。これについては、インターネットメディア協会ができる以前に、例えば先ほど申し上げたJIAAのところで広告表示の問題だとかに取り組んできました、私たちも現段階で特にそういう規定を設けているわけじゃありませんが、当然のことながら、情報の信頼性という意味においては、要するに関係性の明示というのは極めて重要な問題だと思うんですね。

例えば、先ほどおっしゃったようなスポンサーがあるのに表示されていない、あるいは何かを提供してもらっているのに表示されていないということは、記事の信頼性そのものにかかわる問題ですので、当然、関係性は明示されていくことによって信頼を確保していく問題だと僕は考えています。そういう意味で言いますと、私たちが抱える、今、策定中の倫理綱領の中に信頼性の確保ということがありますので、その中で取り扱うべき問題だと思っています。

では、具体的にどういうこととしていくかということに関しては、もちろん何度も言ったような勉強会のようなところでケースを共有していくというのもあるんですが、もう一つは、既にJ I A Aのような知見とネットワークを持っている団体があるので、そういうところと私たちがいろいろな意見交換とか交流することによってより対応を深めていくことが重要だと思っています。特に、実は私個人なことと言うと、コンテンツメディアをやったことがあるので、僕はSmartNewsの前は講談社にいたんですけども、これは新聞でも出版社でもそうですけども、編集と広告というのはファイヤーウォールがあって自立している。これは極めて重要なことなんですよね。でも、極めて重要なことであると同時に、問題認識が共通化されてないという危険もあります。だから、それぞれさっき言ったステマのようなことというのは記者の方は意識されていないわけですよね。ファイヤーウォールがあるからこそ情報遮断されているところがあるので、そういう意味で、例えばJ I A Aが取り組んでいることを、要するに広告側の問題として取り組んでいて、もしかすると例えばこの編集側とか、あるいはこのビジネスモデルをつくっている側がきっちり細部まで認識していない可能性もあると思っています。これはまだ可能性の問題ですけども。

そういう意味で言うと、私たちのインターネットメディア協会という壁のない団体によって、それが要するにJ I A Aのような広告の知見のあるところと連携することによって、例えばコンテンツをつくる側も含めて、そういう広告を含めた関係性の明示の問題とかをちゃんと共通認識をして深めていける。そして、場合によっては何らかの対応に持っていけるような形にすることによって、消費者を守っていけるのではないかというふうに考えています。

【森構成員】 ありがとうございました。

【宋戸座長】 ほかにご質問、ご意見いかがでございましょうか。

では、寺田構成員、お願いします。

【寺田構成員】 成原先生に聞きすればいいのか、どうなのかなと思っているところなんですけど、エコーチェンバーとかフィルターバブルって、気になっているのが、エコー効果、フィルター効果そのものが悪だみたいなニュアンスにどんどんなりつつあって、これって、おそらく適度なエコー効果とか、フィルターの効果というのがあるんだと思っています。でないと、いわゆる仲間、そういったグループ関係であつたりとか、そういったものもつくられないということになりかねないので、こういったことに関してどのあたりが適度なもののかとか、そういった研究とかというのはあるんでしょうか。

【成原准教授】 最近では、エコーチェンバー、フィルターバブルなど関連するさまざまな概念が提起されており、アメリカの法学者、社会学者、心理学者が学際的に議論をしてお

りますけれども、比較的古くからそうした議論をしているのが、先ほども少し名前の挙がったキャス・サンスティーンという法学者です。彼は2000年代の初めごろからインターネット上のグループポラライゼーションを主題化し、似たような考え方を持っている集団が固まって、その中だけでコミュニケーションが閉じ、社会が分極化していくという問題を指摘してきました。

ただ一方で、まさに先生が今おっしゃったように、グループポラライゼーション、集団で固まることにも良い面があるということもサンスティーンは指摘しておりまして、同じような考え方や興味を持つ集団で密接にコミュニケーションを行うことによって議論が磨かれるであるとか、知見がより深まるであるとか、そういうことポジティブな効果もある。他方で、あまり同じ集団だけで固まってコミュニケーションが閉ざされてしまうと、社会全体で多様な見解が流通して双方で議論が行われなくなる。要するに、グループポラライゼーションは必ずしも悪い面だけではなくて良い面もある。ですので、オープンなコミュニケーションと閉じたコミュニケーションとの間で、適度にバランスをとっていく必要があるということは、規範論の観点からも言えるのではないかと思います。それをどのようにプラットフォームの設計などに生かしていくかということは、ご専門の先生のお考えを伺いたいところです。

【宋戸座長】 ほかにいかがでしょうか。今日踏まえたご意見等でもよろしゅうございますが、いかがでしょうか。よろしいでしょうか。

それでは、やや予定の時間よりも若干早いですが、いつもこの研究会は時間をいっぱいいっぱい使った上にロスタイムまで使うことが多いのですけれども、今日はそこまで行かなかったですが、4人のご発表の先生方からは、大変貴重なインプットをいただきました。ありがとうございました。

私のほうから少しだけ感想めいたことを申し上げさせていただきますと、10年ぐらい前に総務省におきましてICT分野における国民の権利保障フォーラムということで、放送が中心でしたけれども、通信も含めて、新しい時代の表現の自由をどうやって確保していくかということについてかなり大がかりな議論をしたことがございます。そのときには、表現の自由を確保するための砦というのは、情報の発信者もそうですし、情報の発信者の団体の中でのいわば自主的な規律もそうです。それから、当然、読者・視聴者といったような人も含めて重層的につくっていくんだ。それはまさに、表現の自由の行使こそが、あるいは相互の批判や論評を通じて表現の自由を確保していくんだというとりまとめを、当時、濱田純一座長のところでされた記憶がございます。

現在、そこから10年たってということでございますけれども、その基本はやはり変わら

ないだろう、基本的な哲学は変わらないだろうと思いますが、その前提となる環境として、大きくその我々の情報流通をプラットフォームサービスが担うようになってきている。あるいはインターネットの役割が大きくなり、伝統的なメディアがビジネスモデルを含めてさまざまな挑戦にさらされている。他方で、新しいメディア、新しいジャーナリズムというものが胎動しつつある。そういった環境の中で、ここでどういうふうに表現の自由、あるいは情報流通の自由、知る権利を確保しながらきちんとした議論をしていくかということがこの研究会の大きな課題であるというふうに、今日、ご発表とご意見を伺っていて、私は感じたところでございます。

そうした中で、特に森先生がご指摘になったことにかかわると思いますけれども、ICT、あるいはプラットフォーム特有の形でのフィルターバブルでありましたり、ターゲティング広告の問題につきまして、この研究会ではこれまでヒアリングをさせていただきましたけれども、プラットフォーム事業者の方々が透明性を確保していただくとか、あるいは一定の措置をとるときに、そのようなプラットフォーム事業者の方々の信頼性、トラストをどうやってつくっていくか。事業者の方々ご自身もそうですし、また、それが外から見える形でどうやって確保できるか。特に、今日ご発表いただきました乾先生、笹原先生からのご知見というのは、それを外からどうやって検証するか、あるいはそれが正しいものとなるにはどうしたらいいかということについて、非常に貴重なご示唆をいただいたと思います。

他方、情報発信者側の部分で言いますと、フェイクニュースを止めるというよりは、全体として言論には言論で、その表現の質を高めていく。この今の状況の中で高めていくということによって、いわばフェイクニュースを退治していくといえますか、良貨が悪貨を駆逐していくように頑張ってください、そういった意味でのトラストをつくっていただくというための取組について、特に瀬尾代表理事から報告いただきましたし、それを支える哲学、あるいはそういった問題の状況について成原先生からご指摘をいただいたものと思っております。また、その中ではとりわけ瀬尾さんからは、政府が果たすべき役割の問題、あるいはリテラシーの向上についての社会や教育の取組についても具体的なお提言をいただきましたので、これらを踏まえて、さらにこのワーキンググループで議論させていただければと思います。

何か追加が特になければ、本日の意見交換はここまでとさせていただきたいと思います。

最後にトラストサービス検討ワーキンググループがこの研究会の下に置かれ、手塚構成員を主査として、さらに宮内先生に大変ご尽力いただいているところですが、その宮内構成員からご報告があるというふうに伺っておりますので、よろしくお願いいたします。

【宮内構成員】 トラストサービス検討ワーキンググループの主査代理を務めております宮内でございます。本日は手塚主査がご欠席なので、私のほうからご報告させていただきたいと思います。

このワーキンググループは今年の1月から始まりまして、既にもう8回の会合を行っています。かなり密な議論を進めておるところであります。トラストサービスといいますと社会基盤となるようなサービスですけども、こういったトラストサービスにつきまして、我が国における現状と課題を整理いたしまして、いろいろな課題をどう解決するか、そういう方策について検討を行っている次第でございます。

今週月曜日、6月24日に開催されました第8回の会合におきまして、これまでの事業者ヒアリングや構成員の意見等を整理した中間とりまとめ案というものにつきまして審議を行いました。中間とりまとめ案では、さまざまなトラストサービスにつきまして制度化に向けてどう検討を進めていくか、制度化をターゲットにして検討を進定めていくといった取組の方向性を示しております。

主に扱っているトラストサービスは、リモート署名、タイムスタンプ、eシールというものでございます。リモート署名というのは、署名のための秘密の鍵をサーバーに預けておいて、これをクラウドを介して使うということで非常に利便性が高まるといった署名でございます。タイムスタンプは、皆さんご存じだと思いますけれども、ある電子文書の、ある時刻に存在したということが証明できるものであります。それからeシールというのは、自然人でなく組織に対する認証、あるいは電子署名、こういうものを実現していくもので、これはいづれにつきましても、ヨーロッパでは非常に進んでいるものですので、これの社会的な意味も含めて、我々の取組の方向性というものを示している次第でございます。

このとりまとめ案につきましては、明日6月28日に報道発表した上で、明後日の29日から7月18日まで、意見公募手続きにかけるとなっております。この中間とりまとめの詳細につきましては、こういったその意見公募手続きの結果も踏まえた上で、次回の本研究会でご報告させていただきたいと思っております。そのときは手塚主査がちゃんと来てくださるものと、私は信じております。

以上でございます。

【宍戸座長】 ありがとうございます。トラストサービス検討ワーキンググループにつきましては、参加されている方々、それから事務局のほうにおかれましても大変ご努力いただいているということで、私からも御礼申し上げたいと思います。それと同時に、楽しみにしております。

【岡本消費者行政第二課企画官】 次回会合につきましては、別途、事務局からご案内いたします。

事務局からは以上でございます。

【宍戸座長】 ありがとうございました。これにて本日の議事は全て終了いたしました。

以上で、プラットフォームサービスに関する研究会第10回会合を終了とさせていただきます。お忙しい中ご出席いただきありがとうございました。