

「デジタル空間における情報流通の健全性確保の在り方に関する検討会」(第14回)・

ワーキンググループ(第10回)第1部

1 日時 令和6年3月27日(水)10時00分～11時30分

2 場所 オンライン開催

3 出席者

(1) 構成員

矢野座長、生貝構成員、石井構成員、越前構成員、奥村構成員、落合構成員、
クロサカ構成員、澁谷構成員、田中構成員、水谷構成員、森構成員、安野構成員、
山口構成員、山本(健)構成員、山本(龍)構成員

(2) オブザーバー団体

一般社団法人安心ネットづくり促進協議会、一般社団法人新経済連盟、一般社団法人セーフ
ティーインターネット協会、一般社団法人ソーシャルメディア利用環境整備機構、一般社団法人
デジタル広告品質認証機構、一般社団法人テレコムサービス協会、一般社団法人電気通信
事業者協会、一般社団法人日本インターネットプロバイダー協会、一般社団法人日本ケーブ
ルテレビ連盟、一般社団法人日本新聞協会、日本放送協会、一般社団法人MyData Japan、一
般財団法人マルチメディア振興センター

(3) オブザーバー省庁

内閣官房、内閣府、警察庁、消費者庁、デジタル庁、文部科学省、経済産業省

(4) 総務省

湯本大臣官房総括審議官、西泉大臣官房審議官、田邊情報通信政策課長、
大澤情報流通振興課長、恩賀情報流通適正化推進室長、内藤情報流通適正化推進室課長補佐、
上原情報流通適正化推進室専門職

(5) ヒアリング関係者

Google Chenie氏、J. J. サヘル氏、フィリン・タン氏、野田氏

4 議事

- (1) 関係者からのヒアリング
- (2) その他

【宋戸座長】 それでは、定刻でございますので、デジタル空間における情報流通の健全性確保の在り方に関する検討会第14回の会合、重ねまして、ワーキンググループ第10回会合の合同会合を開催させていただきます。

本日も御多忙のところ本会合に御出席をいただき、誠にありがとうございます。

議事に入ります前に、事務局より連絡事項の説明をお願いいたします。

【高橋係長】 事務局でございます。

まず、本日の会議は公開とさせていただきますので、その点御了承ください。

次に、事務局より、ウェブ会議による開催上の注意事項について案内いたします。本日の会議につきましては、構成員及び傍聴はウェブ会議システムにて実施させていただいております。本日の会合の傍聴につきましては、ウェブ会議システムによる音声及び資料投影のみでの傍聴とさせていただきます。事務局において、傍聴者は発言ができない設定とさせていただきますので、音声設定を変更しないようお願いいたします。

本日午前の資料は、本体資料として資料14-1-1から14-1-3までの3点、御用意をしております。万が一、お手元に届いていない場合がございますら事務局までお申しつけください。傍聴の方につきましては、本検討会のホームページ上に資料が公開されておりますので、そちらから閲覧ください。

また、ヒアリングシート回答にはURLが記載されているものもございますので、御参加の皆様におかれましては、適宜アクセスしながら御確認ください。

なお、本日は、江間構成員、後藤構成員、曾我部構成員、脇浜構成員は御欠席予定、落合構成員は会議途中での御出席予定、澁谷構成員、山本健人構成員は会議途中での御退出予定と伺っております。

最後に、本日の会議につきまして、報道関係者より、冒頭カメラ撮りの希望がございましたので、構成員の皆様におかれましては、差し支えない範囲でカメラをオンにさせていただきますようお願いいたします。ありがとうございます。

それでは、開始させていただきます。

(カメラ撮り)

【高橋係長】 皆様ありがとうございました。こちらでカメラ撮りを終了いたします。これ以降の撮影は御遠慮ください。

事務局からは以上です。

【宋戸座長】

ありがとうございました。カメラ撮りに御協力いただき、構成員の皆さまありがとうございます。

それでは、本日の議事について御説明をさしあげます。まず、関係者からのヒアリングといたしまして、プラットフォーム事業者でありますG o o g l e様からの御発表、それから質疑を予定しております。早速議事に入らせていただきます。まずは、G o o g l e・C h e n i e様より25分間で御説明をいただき、その後、60分の質疑の時間を設けたいと思います。通訳を入れていただくと伺っております。なお、5分前、1分前には事務局よりアナウンスをさせていただきますので、大変恐縮でございますが、時間厳守で御発表をいただきますようお願いをいたします。それでは、C h e n i e様、どうぞよろしくお願いいたします。

【G o o g l e (C h e n i e氏)】

総務省のスタディーグループ・ワーキンググループのメンバーの皆様、本日は、生成A I時代のコンテンツセーフティーに関するG o o g l eのアプローチ、そして、G o o g l eプロダクトにおける情報の質を確保するための取組について、お話しをする機会を頂戴し、ありがとうございます。

本日は、私、チェニー・ユンがプレゼンテーションさせていただきます。アジア太平洋のコンテンツレギュレーションのリードを務めており、シンガポールのほうから参加しております。また、情報の質に関しては、J. J. サヘル、彼は、コンテンツポリシーに関しても以前皆様とお話しさせていただいております。また、ユーチューブから、フィリン・タン、野田由比子さんも参加しています。

本日のプレゼンテーションは2つに分けております。まず、G o o g l eのアプローチとしまして、我々のプリンシパル、原則に基づいて、こういったアプローチを取っているのか、質の高い信頼できる情報を表示できるような取組についてお話しをしたいと思います。そして、第2部としましては、直近の検知、そして合成メディアの検出と来歴情報の技術に関してお話しをしたいと思います。

私たちは、大胆に物事を進めていきますが、十分に責任を持って進めていきます。これはアルファベット、そしてG o o g l eのCEO、スンダー・ピチャイの言葉です。A Iを責任持って発展させるということは、そのポジティブのインパクトを最大化しつつ、全体的なリスクに対処する、その両方のバランスを取ることが必要となります。

これはかなり繊細な難しいダンスに見えるかもしれませんが、長期的な成功を収めるた

めには、この緊張を私たちは受け入れなくてはなりません。最初から責任を優先することによってのみ社会のウェルビーイングを損なうことなく、A I の持つ変革力を真に発揮させることができると確信しています。

新しいA I プロダクトに着手する前に、まず最初にお伝えしておきたいのは、全ての私たちの業務というのは明確な一連の原則、A I 原則を含めて原則に導かれているということです。2018年、G o o g l e は7つのA I 原則を公表しました。このような原則を公表する最初の企業の一つです。これらは理論的な概念ではなく、私たちの研究・製品開発の基盤となり、また、ビジネス上の意思決定に影響を与える具体的な基準です。

G o o g l e は、既に20年以上にわたって責任あるA I の開発をしてきました。例えば、検索だったり、マップなど、毎日使うプロダクトの中には実はA I が組み込まれています。例えば、2018年に開始したS m a r t C o m p o s e は、少し文章を入力すると完成文を提案してくれるので、簡単にメールを書くことができます。そして、G o o g l e 翻訳だったり、G o o g l e レンズ、これは、例えばO C R、光学文字認識や機械学習といったA I でまさに構築されています。毎日、より多くの方に役立つA I を開発しています。

私たちのアプローチをまとめますと、大胆かつ責任あるアプローチが必要だと考えています。つまり、社会への恩恵を最大化しつつ、その課題に対応しながらA I を開発していくということです。

それでは、我々のポリシーについて、より詳しく触れていきたいと思います。G o o g l e には、生成A I によって作成されたコンテンツを含め、全てのプロダクト、サービスに適用される、長年継続している様々なポリシーといったものがあります。例えば、G o o g l e 広告の不実表示ポリシーの一環としては、操作されたメディア、ディープフェイク、その他ユーザーを欺いたり、搾取したり、誤解させることを意図した偽装コンテンツの使用を禁止しています。

また、この問題が選挙といった非常に重要なタイミングにおいてどう影響するかについても検討しています。例えば、最近、選挙広告に関するポリシーを更新しました。選挙広告にデジタル的に改変または生成された素材が含まれている場合には、広告主がそれを開示することを義務づけています。

私たちは、クリエイターだったり、ユーザー、視聴者、アーティストを含むコミュニティから、テクノロジーが持ち得る影響について常にフィードバックを聞いてきました。そこで、今後数か月のうちに、ユーチューブ上のA I が作成した、または顔や声を含む個人を特

定できるような合成・改変されたコンテンツの削除を、プライバシーリクエストのプロセスを使ってリクエストできるようにします。

また、G o o g l e ディープマインドとG o o g l e クラウドは、共同でA I が生成した画像に電子透かしを入れて識別するための実験的なツールでありますS y n t h I Dを発表しました。この2枚の写真というのは人間の目から見ると全く同じに見えるんですけども、実はデジタルな透かしというのを画像のピクセルに直接埋め込んでおり、人の目では見えないけれども、A I 生成データだと識別できるようになります。

また、ポリシーやツールだけでなく、合成メディア検出のテクノロジーに関しても積極的に研究開発をしています。合成メディアの正確な検出は、生成A I モデルの悪用から保護するための重要なステップの一つです。G o o g l e はA u d i o L M と呼ばれる音声モデルを実験的に使用しています。このモデルは、音声から音楽など、短いサンプルから説得力のある音声を生成することができます。

このA u d i o L M の研究というのはあくまで研究目的でありまして、現時点で広く公開する予定はありません。しかし、この研究の一環として、A u d i o L M によって生成された合成音声を非常に高い精度、98.6%で検出できるクラシファイア、分類器の開発に成功しました。つまり、A u d i o L M によって生成された音声は、聞き手の一部にはほとんど区別がつかないにもかかわらず、テクノロジーを使うと検出が非常に容易であることを示しています。これはA u d i o L M の潜在的な不正利用から保護するための重要な第一歩です。

もう一つのツールとして、アバウトジスイメージというものがあります。これは23年の10月にローンチしたもので、オンラインで見たそのコンテンツを評価するためのツールで、より多くのコンテキストを得ることができます。右のこちらの画像を御覧いただきますと、これはムーンランディング、月面着陸と検索したときのイメージです。これをクリックしますと、過去のヒストリーだったり、どう使われているのかとかメタデータといった情報を得ることができます。また、これがA I で改変されているのかといった情報を見ることも可能です。

25分という制限がありますので、それでは、ユーチューブのスライドに移りたいと思います。ユーチューブでは、改変・合成されたメディアの開示要件とラベル表示を通じてユーザーにコンテキストを提供しています。最近、クリエイタースタジオに新しいツールを導入しました。生成A I を含む改変・合成メディアを使用して実物のように見えるコンテンツを

作成した場合に、視聴者に開示することをクリエイターに義務づけています。この開示というのは、拡大の説明欄のところ、ないしは健康・選挙・金融・ニュースなどデリケートなトピックについては、より目立つように動画自体にラベル表示されることになります。また、ユーチューブがローカライズされている全ての言語、日本語を含めて御利用いただけます。

また、ユーザーコンテキストを高める取組の一環として、メディアリテラシーの向上にも我々は投資をしております。ユーザーが情報を批判的に評価する力を養うための支援として、例えば23年4月、ユーチューブは総務省、国際大学GLOCOMと共同で、偽・誤情報への取組の一環として、「ほんとかな？が、あなたを守る。」キャンペーンを実施しました。このキャンペーンはフェイクニュースが日常生活に潜んでいるということを認識し、情報との付き合い方を考えるきっかけをつくることを目的にしたものです。

それでは、次に、AIの責任に関して適切なバランスを得るためには、複数の視点を持ち寄る協力、コラボレーションが非常に重要ということで、パートナーシップについてお話しをしたいと思います。私たちは、主要な企業とともに責任あるイノベーションへのコミットメントを支持しております。MLCommons、パートナーシップ、オンAI、フロンティアモデルフォーラムといった、業界の主要な取組の創設メンバーでもあります。

また、長年にわたり責任あるAI開発のための基準やガイダンスの研究を行っている研究者や学者のコミュニティを構築してきました。こちらのスライドにあるような学術機関とも連携を取っております。

また、最近の展開としては、2024年2月にC2PA、デジタルコンテンツの来歴記録情報技術に関する国際標準化団体に加盟いたしました。このことを通じて、来歴情報、そして標準仕様の開発に関して、私たちも取り組んでいきたいと考えております。

それでは、少しテクノロジーについてお話しをしたいと思います。まず、安全ということに関して考えてみたいと思います。端的に言えば、私たちはユーザーや社会全体に害を及ぼす可能性のあるもの全てを安全に、というふうに考えております。それから、ディープフェイクに関しましても定義を持っております。

さて、ディープフェイクですけれども、私たちは長年ディープフェイクについて考えてまいりました。そして本日参加されている皆様もよく御存じのとおり、AIを使わないようなメディア操作というのも、写真の編集技術が開発されて以来、ずっと長年存在しているものです。

それでは、技術的な取組として、いわゆる機械がある程度の規模でできることとしてはど

んなものがあるのか。そして来歴ですね。この言葉を本日、何度か使いますけれども、コンテンツがどのように作成されて編集されたか、この来歴についてお話しをしたいと思います。

まず、電子透かし、ウォーターマークです。ウォーターマークとは、メディアコンテンツ自体の一部として、例えば画像やビデオピクセルの中に、通常、生成時や編集時に情報を追加することを意味しています。

次に、フィンガープリンティングですけれども、これはコンテンツのハッシュ値ベクトル表現を作成することです。人間の指紋同様ですけれども、実際、このフィンガープリントにコンテンツの中身というのは一切含まれませんが、ほかのコンテンツのユニークなフィンガープリントと比較して同一であるかどうかを判断することができます。つまり、これはコンテンツがデータベースに保存された生成または認証されたコンテンツのフィンガープリントと一致するかどうかをチェックするために使用することができます。

次に、簡単にメタデータに関して。メタデータとはデータに関するデータという意味ですので、ファイルに添付されている情報で、そのファイルがいつ作成されたのか、いつ変更されたのか。サムネイルなど、そのファイルに関する詳細情報を意味します。そして、機械学習のモデルにおいては、1つまたは複数の生成AIシステムの出力を認識するよう訓練されます。実際には、その学習をした分類器を使用してコンテンツの来歴などを評価することが可能です。

それでは、これだけいろいろなテクノロジーがあるのに、何で問題がまだ解決されないのかですけれども、まず、このテクノロジーというのは全てまだ進化途上にあります。そして、Googleとしては、安全なAIに強くコミットをしており、7つの原則に基づいて、これからもこれらのツールの開発を進めていきますが、それぞれのテクノロジーは完璧ではありません。例えば、メタデータは改ざんされる可能性がありますし、またウォーターマークやフィンガープリントは、悪意のあるアクターであれば損なう可能性があります。また、分類器の精度なども変動する可能性があります。ギャップがあったり、もしくは古くなるおそれもあります。また、これらのソリューションの中には大規模な計算能力を必要とするものもあります。

なので、1つ、こういった万能薬というのはまだ存在しないということを御理解いただいた上で、アプローチを今後組み合わせて、さらに進化していくことをお約束いたしたいと思います。

御清聴いただき誠にありがとうございました。

それでは、最後のスライド、ここからQ&Aに移りたいと思いますけれども、この大変重要なトピックに関して我々の説明をさせていただく機会を頂戴して、誠にありがとうございます。また、事前に、一部の質問に関しては、セキュリティー上、ないしは戦略上、回答できないものもあるかもしれませんが、ぜひ御理解をお願いしたいと思います。

【宋戸座長】 Chenieさん、大変、分かりやすいプレゼンテーション、どうもありがとうございました。AIの研究開発をリードされるGoogle社が、大変優れたアプローチと、それからテクノロジーの開発に取り組んでおられることに共感を覚えました。また、本日、このような総務省のスタディーグループにおいでいただき、丁寧な御説明いただくとともに、この後、可能な範囲でQ&Aに対応していただくということにも、座長として感謝を申し上げます。

それでは、構成員の皆様から御質問、御意見がある方は、私のほうにチャット欄でお知らせいただければとお願いをしたいと思います。ただいまお話ありましたように、お答えいただけるもの、いただけないもの、あるいはお持ち帰りいただいてお答えいただくものも当然あると思いますけれども、この研究会の趣旨でありますデジタル空間における情報流通の健全性の確保という観点から見て必要と思われることを、できるだけ多く、したがいまして、できるだけ簡潔に趣旨を明確に御発言をいただければと思います。構成員の皆様、私にチャット欄で御発言希望をお知らせください。いかがでしょうか。

それでは、奥村構成員。

【奥村構成員】 奥村でございます。どうもありがとうございました。先進的なテクノロジーをたくさんお使いになって多様な取組をされているということがよく分かりました。ありがとうございました。

質問は、主に、ミスインフォメーションやディスインフォメーション対策を、ほかのメディアとか、それから世界的なファクトチェッカーとどのように協力しておやりになるかということについて伺います。例えば、国際ファクトチェックネットワーク、IFCNは2022年の1月にユーチューブにオープンデータを出しております。こちらでは、コンテンツモデレーションポリシーについてファクトチェッカーともっと共有してもらいたいという不満の表明、それから、Googleは削除してしまいますけれども、削除しないでフラグを立てて、それがミスインフォメーションであるということを広く知らせるというような手法のほうの有効ではないかということについて、議論してほしいということなどを

挙げています。そのような、ファクトチェッカーや何かの連携というのが、どういう課題をお持ちなのかということ。

それから、ファクトチェッカーにはもう一つありまして、経済的な問題があります。先ほど業界全体の取組とかパートナーシップというものを御社のポリシーとして挙げられているわけですが、例えばG o o g l e ニュースイニシアチブというようなファクトチェックやなんかでは、御社としては非常に重要な取組だと思われたものについては、かなり予算とか人員とかが削減されております。大きな組織なので御回答は難しいかもしれませんが、リージョンの中ではかなり偉い方が出てきて、今お話しになっていたかと思うので、そのようなレベルで、そのようなポリシーとか考え方みたいなものがどのような理念として共有されているのかということについて伺いたいと思います。ありがとうございました。

【G o o g l e (C h e n i e 氏)】

御質問ありがとうございます。G o o g l e としては、このメディアリテラシー等、非常に重要だとグローバルに考えておりまして、大きな投資もしております。

まず最初に、ジャパンファクトチェックセンターについて申し上げたいと思います。こちらに関しては、セーファーインターネット協会に22年、150万ドルの投資をしており、ファクトチェックの非常に重要な取組の一つだと考えております。具体的なリージョンごとの取組に関しては、私の同僚のフイリンのほうからお願いしたいと思います。

【G o o g l e (H u i L i n 氏)】 皆さん、こんにちは。私のほうからユーチューブについて回答させていただいて、その後J. J. のほうから、より幅広い取組について回答したいと思います。

まず、ファクトチェックに関しての人員の削減だったり、G o o g l e ニュースイニシアチブに関しての人の削減といったお話がありましたけれども、実はこれらの人員だけではなく、より多くの人員がこのファクトチェック等に関わっております。なので、ここで見えている数字だけではない、いろいろな機関などとも、またファクトチェッカーなどとも連携を取ってファクトチェックを行っています。

ユーチューブは基本的にはビデオをホストするところで、視聴者の方はエンターテインメントだったり、スキルの向上だったり、新しいニュースだったりを見るためにいらっしゃいます。なので、ファクトチェックに関しての投資という観点からしますと、例えば様々なビデオフォーマット、ロングタイプ、ショートタイプに関してのビデオファクトチェックを容

易にできるような投資というのもしております。また、昨年6月ですけれども、Googleファクトのオープンセッションとしまして、世界中のファクトチェッカー25を招待して、ファクトチェックのビデオに関してのつくり方だったり、ベストプラクティス、コンテンツなどについての共有を行いました。

【Google (Jean-Jacques氏)】

それでは、私もお話しをしたいと思います。宍戸座長、そして、コミッティーのメンバーの皆様、また皆様とお会いでき、お話ができ、大変うれしく思います。

そして、この3、4年間の対話の中で、皆様よく御理解いただいておりますとおり、私たちは偽・誤情報に関して戦略的なアプローチを取っております。まず、質の高い信頼できる情報を上に上げる、優先的に表示をするということ。そして、偽情報・誤情報であり得るものに関してはフラグをするというのは、もう継続的な取組です。そしてGoogleは、そのためにもファクトチェックの組織だったり、コミュニティをサポートしております。非常にこのエコシステムにおいて重要な役割を果たしていると考えているからです。

日本を特に中心としまして、どんなサポートを提供しているのかについてお話しをしたいと思います。

まず、テクノロジーに関して、私たちはファクトチェックの組織に様々なツールを提供しています。オンライン上の情報のマークアップをより容易にできるように、そしてその情報を受けて、私たち自身も分析を行い、それをユーザーにフィードバックしています。

また、ファクトチェックコミュニティがより体系的に、より能力高く、より大規模に展開できるようなサポートというのもしております。例えば、ポインターインスティテュートにおきましては、1,300万ドルの助成金を22年後半に私たちが支出しまして、65か国80の言語を横断する国際的なファクトチェックの能力を高めるために、様々なケイパビリティ、リサーチなどに支援をしております。

そして、国際的な取組に加えまして、日本に関しまして申し上げますと、Googleの慈善事業部門であるGoogle.orgにおきまして、セーフアーインターネット協会に対して150万ドルの拠出を行い、日本のファクトチェックセンター、ジャパンファクトチェックセンターの立ち上げを支援しました。ここでは、様々なオンラインの情報の研究リサーチを行ったり、そういった誤情報に関してのトレンドだったり、メディアリテラシーを高めるための啓発活動などもしています。

そして、より最近の取組としましては、様々なファクトチェックの組織というのが、特に

アジア太平洋でより連携を取れるように支援をしております。例えば、日本のファクトチェックセンターとインドネシアのファクトチェックセンターだったり、セーファーインターネットラボ、こういったところがより直近のトレンドだったり、ベストプラクティスの情報共有ができるように、そのつながりというのを私たちのほうも支援しております。

ありがとうございました。

【宋戸座長】 ありがとうございました。

この後、さらに8人の構成員から現時点で御質問、御発言の希望があります。今の奥村先生の御質問のように御質問の趣旨を明確にして、それから申し訳ありませんが、1人1問で質問はお願いをいたします。

それでは、続きまして、澁谷構成員、お願いします。

【澁谷構成員】 よろしくお願いいたします。

【宋戸座長】 ちょっと待ってください。Chen i e さんが手を挙げていらっしゃるけれども。

【G o o g l e (Chen i e氏)】 申し訳ありません。

もう一点だけ。4月2日に国際ファクトチェックデイがございますので、私たちのほうもたくさんの取組を予定しております。ぜひそれも御覧いただければと思います。

【宋戸座長】 ありがとうございます。

それでは、お待たせしました。澁谷構成員、お願いします。

【澁谷構成員】 ありがとうございます。よろしくお願いいたします。様々な取組に関しまして、丁寧に御説明くださり、ありがとうございます。

私のほうからは、ユーチューブクリエイターが収益を得る仕組みに関しまして、御質問を申し上げます。ユーチューブクリエイターが収益を得る仕組みに関しまして、実際に収益化の停止・無効化に関するポリシーの日本国内の運用状況に関しまして、差し支えない範囲で、肌感覚でも結構ですので、御教示いただけますと幸いです。例えば、ポリシーに違反したという理由で収益化を停止・無効化した数とか、あるいは具体的な事例などがありましたら、ぜひ教えていただけますと幸いです。

【G o o g l e (Hu i Lin氏)】

それでは、まず、私から、クリエイターがどうユーチューブを通じて収益を得られるかについて簡単に申し上げた上で説明をしたいと思います。

まず、ユーチューブのクリエイターが利益を得るためには、ユーチューブパートナープログラムに加盟しなくてはなりません。このユーチューブパートナープログラムに加盟するためには、1,000人以上のチャンネル登録者数で4,000時間の有効な視聴時間を過去12か月に記録している、ないしは、1,000人以上のチャンネル登録者数で1,000万回以上のショート動画の視聴回数が過去90日になくなくてはなりません。この閾値を満たした場合には、私たちのほうでチャンネルのレビューというのをを行います。つまり、マネタイゼーションを行うためには、コミュニティーのガイドラインだったり、一連のポリシーだったり、著作権に関するポリシー、そういったものを全て遵守していることを確認いたします。そして、YPPに登録されますと、広告がそのビデオでストリーミングされた場合に、そのネットレベニューの55%をクリエイターに支払っています。

一旦、このYPP、ユーチューブパートナープログラムに登録されたチャンネルにおいて、今度はその動画の中に例えば違反するものがある、違反というか、私たちの広告掲載に適したコンテンツのガイドラインに反するものがある場合には、そこで掲載される広告の数がまず減ります。一般的な適した動画に比べますと、クリエイターとしては、より広告が表示される回数が少なくなるので、売上を減らすということになります。例えば、反ワクチンの動画などがこれに当たります。

それから、場合によってはYPPの資格自体を失うということもあります。マネタイゼーションポリシーに違反をした場合、これは広告宣伝からのレベニューシェアを失うというだけではなくて、例えばそのチャンネル登録だったり、ほかのファンディングといったものを得る機会というのも失うということです。

そして、コミュニティーガイドラインに対する繰り返しの違反、ないしは悪質なコンテンツ、例えばポルノグラフィーだったり、児童虐待、性的虐待をする動画、もしくはヘイトスピーチといったものに関しては、チャンネル自体を削除するといった場合もあります。責任を持った対応というのが一番重要です。ユーザーを安全にする、そしてまた同時に、表現の自由を損なわない、そしてクリエイターに対してフェアな対応を取る、こういった形でいろんなバランスを取っております。ありがとうございました。

【宍戸座長】 ありがとうございました。繰り返しになりますが、御質問は簡潔に1問でお願いをいたします。

それでは、山口構成員、お願いします。

【山口構成員】 ありがとうございます。国際大学の山口です。御社の取組について丁寧

に教えていただき、ありがとうございました。また、事前に多くの質問があったかと思いますが、お忙しい中、真摯に取り組んでいただいたことを感謝申し上げます。

私からは、選挙に関して1点、御質問させていただきます。現在、多くの国で選挙を控えておりまして、選挙における偽情報への関心というものが高まっております。御社は選挙についても対策をしております、選挙関連ポリシーがあると思います。こちらの日本国内における運用状況、直近の国政選挙に際しての対応とかについて、差し支えない範囲で、この定量的・定性的なものについて御教示いただければ幸いです。

並びに、特に選挙では、今この生成AIというものが注目されていると思いますが、こういった生成AIによるコンテンツについて、とりわけ選挙のときに何か強めの対策を取る予定があるかといったところについても教えていただけますと幸いです。

私からは以上です。

【G o o g l e (C h e n i e 氏)】 先生、ありがとうございます。

24年はスーパー選挙イヤーということで、私たちも非常に重要だと考えております。今年、世界の半分の人口が大きい選挙を経験します。アメリカ、インド、EUなどにおきまして。そして、G o o g l e もユーチューブも、24時間365日、選挙に関してはサポートを提供しております。先ほどJ. J. からもございましたけれども、重要なのは信頼できる質の高い情報を上に上げること、そして、コンテンツのモデレーションを行って違反している情報、信頼できない情報をダウングレードしていくということです。

また、今は特に生成AIによりまして選挙対策がさらに複雑化しております。不実表示だったり、プライバシー、なりすましといった様々な対策がありますし、先ほど申し上げたとおり、AIの前から様々なディープフェイク等の対策というのは既に行っている。こういったものを全て組み合わせて対策を取っているところであります。

【山口構成員】 ありがとうございます。

【宍戸座長】 ありがとうございます。

それでは、山本健人構成員、お願いします。

【山本（健）構成員】 山本です。

私からは、先月のミュンヘン安全法の会議に関して質問させていただければというふうに思います。そちらでは、選挙におけるAIの悪用との闘いに向けた技術協定で、8つの項目について合意がされたというふうに理解しているのですが、これも先ほどの選挙と関連しますが、それぞれの項目について、特に日本国内で今後選挙が行われる場合に実施を予定

されている対策があれば、差し支えない範囲で教えていただければと思います。また、日本国内に特化したものに限らず、日本にも適用されるグローバルに実施予定の取組でも構いませんので、教えていただければと思います。よろしくお願いいたします。

【G o o g l e (J e a n - J a c q u e s 氏)】 それでは、私のほうから。

まず、この大変重要な取組に関心を持っていただき、ありがとうございます。国際的なレベルで、業界横断的に安全な選挙を行うための一連の取組です。先ほどのC h e n i eの説明に関連しますけれども、選挙の公正さに関しましては、3つの柱があります。

まず、有用な情報を有権者に伝えるということです。例えば、質の高い、信頼できる情報を検索結果の上位に上げるといったものが含まれます。

2つ目としましては、プラットフォームを不正から守るということです。これは一連のポリシーだったり、様々なツールなどを使って安全なプラットフォームを維持するということになります。

そして3つ目、これが特にこのトピックで重要となりますけれども、A Iで生成されているか否かを識別するというものです。私たちは様々なツールやテクノロジーに投資をしております、A Iで生成されたコンテンツか否かというのをより容易に判別できるようにしております。例えば、ユーチューブでは、改変された、ないしは合成された動画であるということをクリエイターがラベルづけ、表示することを求めています。これは、実物のような、実物と間違ふようなコンテンツの場合です。

また、選挙という観点からは、G e m i n iに関しまして、非常に慎重なアプローチを取ろうということで、この数週間、新しいプロトコルを導入しています。これはG e m i n iが回答できる選挙関連のクエリを制限するというものです。しかし、ユーザーを、そして社会を本当に守るためには、社会全体の取組が必要であります。まず最初には、業界横断の取組ということで、このようなパートナーシップに私たちは本当に能動的に積極的に創設に関わっております。

2つ、メインの取組があります。まずはユーザーに知らせるということです。A Iで制止されたコンテンツですと、ユーザーにオンラインで知らせるというものです。もう一つが、C 2 P A、コンテンツ来歴及び信頼性のための標準化団体であります。これは業界連携で透明性を高める。そして来歴だったり、標準化コンテキストを高めていくというものです。具体的には、なるべく多くのプラットフォームを巻き込んで、情報の来歴に関してのラベルづけというのを行っていく予定です。

そして、先ほどおっしゃっていただいたテックアコードですけれども、こちらもMe t a、M i c r o s o f t、我々といった世界のリーディングテクノロジー企業が約束をして、A Iの動画、イメージ、オーディオなどを用いて、不正に選挙等に介入を行わないようにする、そのために業界としてテクノロジーを展開するというものです。

そして、このテックアコードだけでなく、その他の取組も含めまして、私たちは、ツールを様々な形で共有していく、そしてまた、トレンド情報などの情報の共有も行っていく。そのことで、悪意あるアクターがA Iを使って選挙に不正に介入しようとして、偽情報・誤情報を打ち出すのを防止するということを取り組んでおります。例えば、ディープフェイクのデータセットを私たちは共有します。このデータセットを使えば、ディープフェイクをより容易に他社も検知することができ、責任あるA Iに向かって進むことができます。

また、不正コンテンツに関しても、例えばハッシュ値などを業界の様々なアクターと共有していきます。そして、これは悪いコンテンツですよといったことを他社も簡単に検知できるようにする。私たちのそれぞれが持っているリソースを集めて、効果的に悪意のあるアクターと闘いたいと考えています。彼らは、別にどのプラットフォームでもいいわけですから。

【G o o g l e (H u i Lin氏)】 J. J. に加えまして、山本先生、山口先生の質問の中で、ポリシーに違反をした際、どうその違反を検知しているのかといった側面もあったかと思しますので、少し申し上げたいと思います。スタディーグループの方はよく御存じのとおり、我々はG o o g l eのスレッドアナリシスグループとも連携を取っておりまして、国内・国外のコーディネートされた選挙への介入に対しての情報というのを集めております。

また、このTAG、脅威分析グループを通じて、業界の他社ともベストプラクティスの共有をしておりますし、また法執行機関などとも情報共有をしております。また、TAG Bulletinを通じた情報共有を行い、例えばこれは悪いアクターだということを特定したら、より簡単に削除することが可能になります。

また、ユーチューブとしましても、こういったリスクを判定するためのチームがおりまして、ユーザーエクスペリエンスだったり、コミュニティにとって有害であるといったものを分析・研究しております。特に、最近では生成A Iに関して特化しています。

また、選挙に関しましても、その選挙に特化したコンテンツモデレーターがモデレーションを行っておりますし、また、コミュニティガイドラインなどのほうも、その選挙に関連したところというのを見ております。また、ローカルランゲージでのサポートというのも行

っております。

私たちは、プラットフォームを、その重要な選挙のタイミングにおいてもぜひ安全に保ちたい、そのために闘っていきたいと考えております。また、本日お話しをした内容のほとんどが人のモデレーションに関してでしたけれども、1分当たり500時間の動画がアップロードされているという、この規模を考えますと、非常に機械学習も重要で、マシンラーニングにも投資をしております。また、決して我々単独ではできませんので、政府だったり、市民社会のサポートにも感謝しており、今後もぜひこの悪質な選挙への介入を防止するために、総務省様、そしてスタディーグループの皆様、市民社会と連携を取っていきたいと思います。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。

まだまだ御質問がありますので、簡潔に御質問いただき、また御回答いただければと思います。それから、次のクロサカ構成員からの質問は、英語で行われます。そして、それについては、クロサカ構成員自身が簡潔に日本語で構成員にも説明します。その後、Google様から御回答いただくようお願いいたします。

【クロサカ構成員】 （英語発言）電子透かしやフィンガープリント技術というのは、限られたコミュニティー、つまり、ユーザーやデータが予め識別・認識されているコミュニティーで非常に有効に働く。C2PAも同じように見えるものなのですが、一方で、インターネットの空間というのは、エンドユーザーやデータというのは全て認識し切ることは難しいという特性もあります。今後これをどのようにより幅広く、識別されていない世界に拡張していくようにお考えかということをお伺いしたものです。よろしくお願いいたします。

【Google（Chen氏）】 クロサカ先生、ありがとうございます。

まず、私たちとしても取組が始まったばかりです。C2PAに私たちがサインしたのは24年2月という、まだ直近であります。また、私の最後のスライドにありましたとおり、非常に複雑な問題で、1つの万能薬はなく、複数のアプローチやテクノロジーを組み合わせで対策をしていくことが必要だと考えております。そしてまた、メタデータに関しても剥がすことが可能ですし、改ざんも可能です。ウォーターマークやフィンガープリントに関しても、並べかえたり編集をしたりといった形で損なわれるおそれがあります。いただいた懸念というのは非常に重要なものですし、今後、私たちとしましても取組を進めていきたいと考えておりますが、まだ始めたばかりというところです。

同僚のほうから何かコメントありますでしょうか。

【G o o g l e (H u i Lin氏)】 ユーチューブとして申し上げたいと思います。私たちはバランスについていろいろ考えております。1つクリアなのは、生成A Iが決して悪いものではないということです。クリエイターはクリエイティブなプロセスでいろいろな形で生成A Iを使っています。非常にパワフルなツールでクリエイティブなプロセスをサポートするものです。より多くのユーザーがクリエーションを行ってグローバルプラットフォームにアクセスできるものと考えておりますので、ラベル表示のポリシーに当たっても、例えばその生産性を上げるために生成A Iを使っているだけであれば、別に表示をする必要はありません。例えば、スクリプトの作成だったり、アイデア出しだったり、字幕の自動作成、そういったものに生成A Iを使っても別にラベル表示は不要ですし、また、合成されたコンテンツがどう見ても現実的ではない、実物的ではないのであればラベル表示は不要です。

こういったアプローチを使って、ユーザーとクリエイターの間での信頼関係をより強化したエコシステムにつなげていきたいと考えております。ありがとうございました。

【クロサカ構成員】 ありがとうございました。(英語発言)

【宋戸座長】 それでは、生貝構成員、お願いします。

【生貝構成員】 大変貴重な御説明ありがとうございました。

御案内のとおり、E Uではデジタルサービスアクトに関連して、偽情報の行動規範、コード・オブ・プラクティスがつくられていたり、あるいは、昨日はちょうど選挙プロセスのインテグリティに関わるガイドラインが欧州委員会から公表されたところであります。そうした中で、我が国でも、日本でもそうしたガイドラインやコード・オブ・プラクティスをつくっていくことの必要性・妥当性について、もしお考えがG o o g l e様としてあれば、あるいはそういうものをつくっていく中で、特にこうした点に気をつける必要があるのではないかということについて、もしお考えがあれば教えていただければ幸いです。よろしくお願いします。

【G o o g l e (C h e n i e氏)】 ありがとうございます。

D S Aについての御質問ありがとうございます。E Uにおきましては、我々G o o g l eのような組織に対しても、コード・オブ・プラクティス、行動規範といったものも示しておりますし、リスク軽減のための取組だと考えられますが、必ずしもこのアプローチというのが全ての規制当局に対して適切なものかに関して、私たちとしては言及する立場にはないと考えておりますが、J e a n - J a c q u e sのほうから何かございますでしょうか。

【G o o g l e (J e a n - J a c q u e s 氏)】 もちろん、ほかの法域で何が起きているのかを調査するのは重要なことですし、EUに関しても長年の取組を経て、このデジタルサービス法に至ったということです。また、世界中で様々なコードといったものが存在しており、業界主導型のものもあれば、EUのように政府主導型のものもあります。ただ、私たちとしては、何がその国にとってベストなのかというのは、その国次第だと考えております。

それを背景としまして、日本のステークホルダーとの私たちの取組について少しお話しをしたいと思います。ここ二、三年、様々なインタラクションだったり、協働の取組というのを、例えばソーシャルメディア利用環境整備機構だったり、セーフティーインターネット協会とともに重ねてまいりました。意識啓発といったユーザー向けの取組などをいろいろ行ってきたんですけども、今後、私たちが日本でぜひスケールアップしていきたいと考えているのは、例えばテックアコードのような取組を日本において具体的に展開していくということです。偽情報・誤情報と闘うためのツールだったり、データセットだったり、情報の共有というのを、役に立つ形で実践的な形で、業界の様々なプレーヤーだったり、ファクトチェックの組織などと共有していく、こういった取組を具体的に進めるというのが、今、私たちが一番考えていることです。

【生貝構成員】 ありがとうございます。

【宋戸座長】 ありがとうございます。

越前構成員、お願いします。

【越前構成員】 国立情報学研究所の越前でございます。大変興味深い発表、どうもありがとうございます。

私からは、ユーチューブにおけるポリシーについて御質問させていただきたいと思います。御社は、ユーチューブにおいて、AIによって生成されたコンテンツについてのラベルづけ機能の実装や義務化を発表しておりますけど、悪意のある人はラベルづけをしない懸念があります。そこで以下を質問させていただきます。

今後、サービス上の分析やメタデータの活用などで、例えばシンセティックメディアディテクションなどの技術的手段を用いて自動でラベルづけがされるようにする予定はあるでしょうか。また、ない場合は何が問題になっているのでしょうか。

また、さらに、ユーチューブのポリシー上、ラベルづけがされなかった場合の対応として、具体的にどのような内容、例えばそのコンテンツの削除とか、ラベルの強制的な付与などを

想定されているのかを御教示ください。

最後に、上記のユーチューブのポリシーにおける日本国内における運用状況、例えば違反の認知、執行件数だとか、違反の存在を外部から指摘された件数と、それに応じてコンテンツを削除等した件数などを差し支えない範囲で御教示ください。

以上です。

【G o o g l e (H u i Lin氏)】 御質問ありがとうございます。

まず、まだまだ初期段階だということを申し上げたいと思います。ファーストフェーズを導入したのが2月で、私たちの生成A I ツールでありますドリームトラック、ドリームスクリーンといったユーチューブのA I ツールを使ったものに関しては、自動的にラベルを付与するといったものを開始しております。

先ほど申し上げたラベルづけに関しての第2ステップというのは、実はつい先週、クリエイタースタジオのほうで提供開始したものです。合成データもしくは改変されたデータが実物に近い場合はラベルづけするというのは、まだ先週導入したばかりのものであります。今後、必ずしも全てのクリエイターがこのラベルづけをしないということはございますので、継続的に常にそういったラベルづけをしないようなユーザーに関しては、例えばそのコンテンツの削除だったり、ユーチューブパートナープログラムの資格停止だったり、その他ペナルティーといったものも検討してまいります。

ここで強調しておきたいのは、こういったペナルティーというのはまだ一切科しておりません。まだまだ導入したばかりですので、クリエイターの方々にも、ある程度経験してもらって新しいプロセスに慣れていただく時間が必要だと考えております。また、生成A I を使うのが悪いことといったメッセージは出したくないと考えています。クリエイターが自主的に開示することを期待していますけれども、場合によっては、ラベルづけされるべきコンテンツといった形で強制的にラベルをつけるといったことも行います。これは特にセンシティブなトピック、選挙だったり、ニュース、金融、それから健康などにまつわるコンテンツの場合です。

まだテクノロジーは完璧ではありませんが、日進月歩で検知のテクノロジーというのは進化しています。クリエイターが合成ないしは改変されたコンテンツであることをディスクローズするように、テクノロジーのほうにも力を入れていきます。

削除されたコンテンツ等に関しましては、日本では例えば23年のQ3においては、コミュニティガイドラインに違反した日本IPの動画として10万動画を削除しております。

この10万の削除された動画のうち97%は、機械学習によって自動的にフラグgingがされたものです。また、動画のうち93%は100回以下の視聴ということで、少ない視聴のタイミングで削除されています。これらのマトリックスを見ますと、ある程度私たちは日本のユーザーを安全に守っており、悪いコンテンツも速やかに除去できていると考えております。ただ、完璧ではありませんので、政府や市民社会、様々なパートナーの御支援を今後もぜひいただきたいと考えております。

【宍戸座長】 Chenieさん、大変、分かりやすいプレゼンテーション、どうもありがとうございました。AIの研究開発をリードされるGoogle社が、大変優れたアプローチと、それからテクノロジーの開発に取り組んでおられることに共感を覚えました。また、本日、このような総務省のスタディグループにおいでいただき、丁寧な御説明いただくとともに、この後、可能な範囲でQ&Aに対応していただくということにも、座長として感謝を申し上げます。

それでは、構成員の皆様から御質問、御意見がある方は、私のほうにチャット欄でお知らせいただければとお願いをしたいと思います。ただいまお話ありましたように、お答えいただけるもの、いただけないもの、あるいはお持ち帰りいただいてお答えいただくものも当然あると思いますけれども、この研究会の趣旨でありますデジタル空間における情報流通の健全性の確保という観点から見て必要と思われることを、できるだけ多く、したがって、できるだけ簡潔に趣旨を明確に御発言をいただければと思います。構成員の皆様、私にチャット欄で御発言希望をお知らせください。いかがでしょうか。

それでは、奥村構成員。

森構成員をお願いします。

【森構成員】 御発表ありがとうございました。大変示唆に富む御発表でした。

私から質問1つということですが、2つの情報についてお聞きしたいと思います。

1つ目は、なりすましです。どんなものかということですが、例えば、日本で有名な実業家の画像を使って、これは本物です。AIでつくられたものではありません。その人が、あたかも次のように言っている。私が経済的に成功した理由を皆さんに共有します。有料のセミナーに参加してください。しかし、その人自身はそんなことは言っていないし、そんなセミナーには全く関係がないわけです。典型的にはこのようなものになりすましですが、その人が発信していない情報を、あたかもその人が発信したかのようにしたもの。それについて何か制限をされているでしょうか。このような情報は、広告として出稿されること

もあれば、ユーチューブに投稿されることもあると思います。両方について教えていただければと思います。また、制限する、削除する条件みたいなものがあるのでしょうか。例えば、通報があれば制限する。例えば、財産的損害が生じる可能性がある、詐欺的被害が生じる可能性があるあれば、削除する。そんなことは関係なく、なりすましであるということが分かれば、その人が言っていないということが分かれば削除する。これについて教えていただきたいと思います。

すみません、もう一つありまして、政治的広告です。どんなものをイメージしているかといいますと、選挙に関することではないのです。例えば、特定の国内の人または団体が日本の国益に反することをしているとか、特定の外国が急速に日本にとっての脅威になっているとか、あるいは日本の政府の一部門が政策的に失敗していることによって日本の経済が大きなダメージを受けようとしているというような、そういう受け取る人の政治的信条に働きかけるようなメッセージのようなものが広告として来たときに、それを制限されることはあるでしょうか。もし、制限されているとすれば、何か条件、基準があるでしょうか。例えば、真実でない場合には制限するとか、マス広告であれば出すけれども、行動ターゲティング広告であれば出さないとか、そのようなことがあれば教えていただきたいと思います。

以上です。よろしくお願いします。

【G o o g l e (C h e n i e 氏)】 すみません、確認ですけれども、今、政治的な広告とおっしゃいましたが、選挙に関連したものではないとおっしゃっていました。これはつまり、例えば外国からの選挙介入といったお話についての質問ではないということですね。これに関しては、私たちのほうでも専門のグループがありまして、外国からのサイバー攻撃などに備えていますけれども、そういった話ではないという理解で正しかったですでしょうか。

【森構成員】 その話ではありません。

【宋戸座長】 すみません、私から1回介入させていただきますが、今の御質問の趣旨は、もうちょっと分かりやすく言いますと、選挙に限らない一般的な公共的な事柄に関する世論形成への介入の問題だと御理解いただくのが分かりやすいかと思います。

【G o o g l e (J e a n - J a c q u e s 氏)】 御質問ありがとうございます。御存じのとおり、G o o g l e は、たくさんのポリシーを持っておりまして、例えば有害コンテンツを規制していますけれども、そこには、いわゆる偽情報、そして誤情報に関してのもの、

不実表示に関するものがございます。私たちのプロダクト全般にこれらのポリシーは当てはまりますし、広告も含めてということになりますけれども、例えば危険な広告、もしくは他者の名誉を傷つけるような広告、もしくは選挙だったり民主的なプロセスを毀損するような広告、そしてまた、市民社会に混乱を来すような広告、そういったセンシティブな内容というのも含めますし、また、自然災害といったときの情報といったものも含まれます。

例えば、G o o g l e は選挙運動の一環であるサイトを宣伝する広告を、日本では表示することを許可していません。また、広告主が政治的所属に基づくターゲティング広告を行うことも許可していません。責任のある広告行動という形でポリシーを取っておりますので、例えば政治的なメッセージ、気候変動に関しての不正確な情報といったものもこれには含まれております。

冒頭の質問に戻りますけれども、不実表示のポリシーというのは、G o o g l e の長年のポリシーの一つであります。これはその当事者が指示していないにもかかわらず、当事者と関連している、ないしはその当事者による指示を受けているといったことを誤って偽情報として提供するものです。このようなものに関しましては、広告の削除をいたします。そして繰り返し行われているようであれば、アカウント自体の削除も行います。そして、より強固なポリシーという形で許可されないビジネス手法に関してのポリシーも当てはまります。これは具体的に金銭を奪う詐欺を行う目的で、公的な人物との関連性ないしはその人の指示をうたっているというものです。これに関しては、より強固な対策がとられます。

この許可されないビジネス手法の違反というのは、悪質で深刻なものと捉えられますので、違反していることが判明したGoogle 広告アカウントは、事前の警告なく強制停止され、今後広告を掲載できなくなります。こういったなりすましに関して、特に公的な有名人のなりすましに関しては長年取り組んでおりますし、さらにポリシー等を強化していきたいと考えております。広告だけでなく、ほかのプラットフォームに関しても同様です。

【宍戸座長】 ありがとうございました。

それでは、増田構成員、簡潔に御質問をお願いいたします。

【増田構成員】 ありがとうございます。

利用者の立場からの質問になります。先ほど市民との連携について言及がありましたけれども、近く、個人の顔や声を模倣したコンテンツの報告というシステムを設定すると伺いました。これについては、利用者がそのようなコンテンツに気がついた場合に通報するという仕組みであるというふうに理解してよろしいでしょうか。それについては、既に問題ある

広告の通報ができるように聞いておりますけれども、それとの違いや、それから利用者が使いやすいような方法、何かお考えがあれば教えていただければと思います。

【G o o g l e (C h e n i e氏)】 ありがとうございます。

今いただいた質問というのは、ユーチューブの生成A I を使ったコンテンツのラベルづけのことで理解してよろしいでしょうか。

【増田構成員】 通報をするというシステムなので、利用者のほうから情報提供をするというシステムがあるのかどうかということです。

【宋戸座長】 増田構成員、ユーチューブの話か、それともG o o g l eの別のサービスの話かという点はいかがですか。

【増田構成員】 ユーチューブだけではないです。G o o g l eの検索エンジンのほうでも同じです。

【G o o g l e (C h e n i e氏)】 ウェブサイト上、G o o g l eポリシーはたくさんあるんですけれども、なりすましだったり、プライバシーだったり、詐欺・不正だったり、そういったポリシーの違反がある際には、ウェブサイトから簡単にユーザーの方は通報することが可能です。これは真実でかつ質の高い情報を提供するための取組の一環であります。

ほかに何かございますでしょうか。

【G o o g l e (H u i L i n氏)】 ユーチューブに関してですけれども、ユーチューブのコミュニティーガイドラインの通報を行うに当たっては、ヘルプセンターといったものも設けております。ですので、動画だろうと説明文であろうとチャンネルのラベルであろうと、ユーザーが簡単にフラグgingをできるようにしております。ご質問に関連する、なりすましとか不正ポリシー、詐欺等のポリシーに関しては、例えばアプリ上で簡単に通報するフラグを立てることが可能です。

もう1点、先ほど説明をしました生成A I を使った改変や合成に関しての通報に関しては、今後、プライバシーの苦情、プライバシーに関しての通報を行うためのウェブのフォームというのをより使いやすいように変更する予定です。これは今後数か月かけて、通報しやすいように改善をしていく予定です。

【G o o g l e (C h e n i e氏)】 こういったユーザーがこういった形でアクセスできるのかに関しましては、先月頂いている質問票の回答のほうで、偽情報、選挙情報、医療情報、なりすまし情報、スパム情報などに関する通報の情報というのを含めておりますの

で、そちらも参照いただければと思います。

【宍戸座長】 ありがとうございます。

お約束の時間でございますので、ただ、もう2点御質問があります。これについては、質問はこの場で御紹介をさせていただきますが、御回答は後ほど別の構成員からも追加であると思いますので、それと併せて可能な範囲で御回答いただければと思います。

まず、山本龍彦座長代理から、ユーチューブコミュニティーガイドラインの適用について、国別動画数が掲載されているところ、日本は9万4,418件で、件数で言えば全体の14位とされているということを前提にして質問です。その御質問の内容は、削除率はどのくらいなのか。そして削除率で見たときの日本の順位は、世界各国の中で何位になるのか、削除理由に関して日本に特徴的な傾向があるのかというものになります。

引き続き、落合構成員からの質問が2つございます。

1つは、英国におけるメディアのコンテンツ表示方法に関するプロミネンスの法整備がされている。特に災害発生時などの非常事態において、新聞・放送等の伝統メディアのコンテンツのファクトチェック記事、また信頼性の高い発信元が発信する情報を優先的に表示するといった取組を、もしGoogleがされているようであれば、差し支えない範囲で教えてください。今後の取組の可能性についても教えてくださいというのが第1問です。

第2問は、広告を配信する先の第三者運営メディア、広告媒体について、偽・誤情報を発信・拡散するような悪質なサイトであった場合に、広告主から当該サイトへの広告掲載を停止するよう求める通報窓口のようなものがあるでしょうか。通報後の対応フロー、実際の通報件数、対応した件数などの運用状況について差し支えない範囲で御教示くださいというものです。

これらにつきましては、この場でなくて、すみません、Googleの日本法人の方を改めて、Googleの皆さんにお伝えいただければと思います。

いずれにいたしましても、本日はGoogleの日本の担当の方、アジアの方、それから、この3年いつもお世話になっているSahelさんに来ていただき、非常に有意義なディスカッションを、ポイントを絞ってすることができたと思います。この点感謝申し上げますという、感謝の言葉だけお伝えいただければと思います。

【Google (Chen氏)】 それでは、Googleチームを代表しまして、座長の宍戸先生、そして皆様、本当にお時間頂戴し、ありがとうございました。Googleとしましては、今後も総務省、そして皆様との取組を進めていきたいと思っています。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。

先ほど申し上げましたように、もし追加で構成員の皆様から御質問等がありましたら事務局にお寄せいただき、それを一括してG o o g l e様にお送りして御回答いただくというプロセスを取りたいと思います。

それでは、最後に事務局より、この午前の部について、何か関連して連絡事項ございますでしょうか。

【高橋係長】 ありがとうございます。

本日午後の会合につきましては、14時より実施をいたします。以上でございます。

【宍戸座長】 ありがとうございます。

それでは、以上をもちまして、デジタル空間における情報流通の健全性確保の在り方に関する検討会の第14回会合、併せてワーキンググループ第10回会合同会合を一時閉会いたします。第2部会合は14時から、あと2時間後ということになります。再び御出席をいただければと思います。ひとまず、ここで散会いたします。ありがとうございました。