



資料

スマホ（SNS、ゲーム、インターネット） との上手な付き合い方

サブテーマ②「青少年の生成AIやSNSの上手な使い方について」

令和8年8月4日

神奈川県福祉子どもみらい局青少年課



生成AIに指示を入力すると、文脈に沿った回答が出力される

指示

※プロンプトという

雨の日の旅行の
アイデアを考えて



生成AI

インターネット上の
文書などを大量に学習

回答

雨の日でも楽しめる
旅行のアイデアを
いくつかご提案します。

- 温泉
- 美術館巡り
- 工場見学 ...

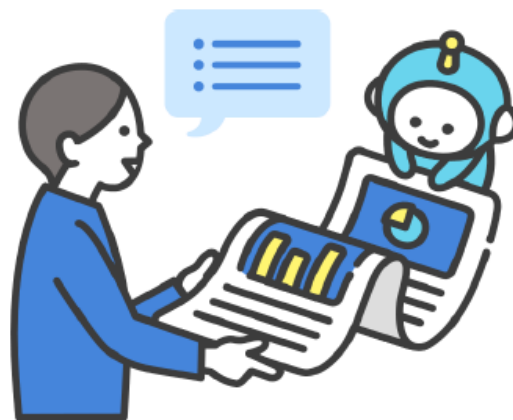
生成AIの特徴

生成物



- 文章・画像・動画 など幅広く作成可能
- 文脈を理解した回答や、人間が思いつかない回答を作ることができる

指示方法



- 高度な技術は必要なく、簡単に使える
- 文章のほか、画像・音声・グラフなどを入力しても回答を得られる

生成AIは幅広い場面で活用できる（以下は代表例）

テキスト生成



文章の作成・要約



情報検索

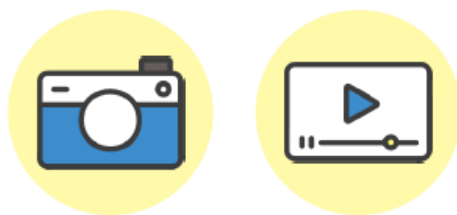


翻訳



議論のパートナー

画像・映像生成



写真・イラスト、
アニメ等の作成

音声生成



音声・音楽の作成

その他



3Dモデルの作成

事例

生成AIにより偽・誤情報が生成される可能性

偽・誤情報の事例 ①

ある生成AIサービスに以下の指示を入力すると、問題のあるリストが生成された。

西日本で最も高い山のTOP10を教えてください



生成されたリストの課題

- ・ 実在しない山の名前が含まれる
- ・ 標高が不正確

XXX駅周辺で
平日ランチにおすすめな和食
のお店を5つ教えて。
予算は1人あたりXXX円以内。
XXXから徒歩5分以内。



生成されたリストの課題

- ・ 存在しないお店が含まれる
- ・ お店の場所が全く違う
- ・ 予算を大幅超過

偽・誤情報の事例 ②

ドローンで撮影された静岡県の水害。
マジで悲惨すぎる...



午前4:39・2022年9月26日・Twitter for Android

2022年9月、台風15号による水害被害が発生している静岡県の画像がSNS上で拡散。その後投稿者は、**画像生成AIで作成した偽画像**だったと公表。

事例

AIが学習した作品に似た画像が生成される場合があるとして、著作権上の論争を呼んでいる

海外では、アーティストが自分たちの作品が画像生成AIの学習に使われ、著作権が侵害されたとして、画像生成AIの開発元に対して訴訟を起こした

ハンバーガーを持つロボットを、〇〇氏の画風で

等を画像生成AIに入力



× Bad

著作権侵害の知識がないユーザーが生成した画像を外部に発信することで、著作権を侵害し、訴訟等に発展する可能性

画像生成AIの出力 ※1



アーティストの作品 ※1



事例

モラル① 学生が自分で考えるべき場面で、安易に生成AIを使ってしまう



あなたは中学2年生です。『走れメロス』の読書感想文を書いてください。



AI

『走れメロス』は宮沢賢治の名作で、忠義と友情の深さを描いた感動的な物語です。私が最も印象的だったのは主人公メロスの人間性とその行動の強さです。メロスの友情と忠誠心は、最も難しい状況でさえも変わらないことを教えてください。

物語の中で、メロスは暴君ディオニスに対して「友情とは何か」を証明するために、厳しい試練を受けます。それは彼の友達セリヌンティウスの命と引き換えに、結婚式に間に合うという約束を果たすというものでした。

私はこの部分が特に心に響きました。メロスは自分自身を危険にさらしながらも、友達への約束を破ることなく、時間厳守で村に戻ってきます。その姿は、どんな困難な状況でも信念を持ち続けることの大切さを教えてください。(中略)

全体として、『走れメロス』は人間の持つ美しい感情と強さを描いた作品で、読んで本当に良かったと感じています。

私自身もメロスのように、友情を大切にし、誠実に生きていきたいと思います。

事例

モラル② 生成AIが出力した偏見のある回答をそのまま使用してしまう

画像生成AIで出力した以下の画像はいずれも**性別**や**人種**に偏りが見られる

※ 現在では、生成AIサービス側の改善が進んでいる

看護師



国会議員



× Bad

複数人の看護師の画像作成を指示したところ、**女性ばかり**が出力され、**人種にも偏り**

× Bad

国会議員の画像作成を指示したところ、**男性ばかり**が出力され、**人種にも偏り**

Note: 画像生成AIで生成。プロンプトとして、左の画像は「Various nurses」、右の画像は「国会議員」と入力

個人情報や機密情報を生成AIに入力すると、情報流出のリスクがある



生成AIは、
利用者が入力したデータを
学習データとして利用
することがある



リスク 1



個人情報^{※1}や社外秘の機密情報^{※2}
を入力すると、他人の質問への回答に
使われ、情報が漏洩する可能性がある

リスク 2



漏洩した情報がサイバー犯罪などに
悪用される恐れがある

※1：個人情報の例：氏名、電話番号、住所、メールアドレス、生年月日、公的ID番号、学校・職場の情報、顔写真など

※2：機密情報の例：経営・営業に関する情報、顧客や社員のデータ、技術情報、法的文章、業務上の秘密、議事録など

生成AIポルノ



芸能人の偽の性的画像 メルカリなどで大量販売 対策求める声

2024年12月14日 16時46分

AIを使って、本物と見分けがつかないほど巧みな偽の画像や動画を作るディープフェイク。この技術を使って、芸能人の顔を合成したとみられる性的な画像がフリマアプリ大手「メルカリ」などで大量に販売されていることが分かりました。

芸能事務所などで作る団体からは、対策の強化を求める声が出ています。

Kanagawa Prefectural Government

約100人 偽の性的画像が...



「メルカリ」では、芸能人の顔を合成した下着姿などの性的な偽の画像が数百円から数千円で販売されています。NHKがことし8月から今月にかけての出品状況を確認したところ、俳優や歌手、アイドルなど少なくともおよそ100人、のべ1000点以上の偽画像が掲載されていました。

画像は「ディープフェイク」と呼ばれるAIを使った技術で作られているとみられ、説明欄に服を脱がす加工をするなどと記載しているものも多くあり、一部は販売済みとなっていました。

また、数は少ないものの、同様の性的な偽画像の販売は▽ネット通販「アマゾン」、▽オークションサイト「ヤフオク!」、▽フリマアプリ「楽天ラクマ」でも確認されました。

出典：NHKホームページ 2024年12月14日

生成AIポルノ



生成AIで作成のわいせつ画像 販売した疑い 4人逮捕 全国初摘発

2025年4月15日 19時45分 生成AI・人工知能

生成AIを使って作成したわいせつな画像をインターネットのオークションサイトで販売したとして、警視庁は20代から50代の4人を逮捕しました。作成されたのは実在しない成人女性の画像で、警視庁によりますと、生成AIで作成したわいせつ物の販売が摘発されるのは全国で初めてだということです。

逮捕されたのは、愛知県に住む小売業の水谷智浩容疑者（44）や、埼玉県に住む会社員の菅沼貴司容疑者（53）ら20代から50代の合わせて4人です。

警視庁によりますと、4人はそれぞれ去年10月、生成AIを使って作成した女性のわいせつな画像をポスターにして、インターネットのオークションサイトで複数回販売した疑いが持たれています。

Kar

無料で提供されている生成AIのソフトに、ネット上にある大量の画像を学習させたうえで、「脚を開く」などとポーズを指示することで、実在しない成人女性の裸に見える画像を作成していたとみられています。

ポスターは、「AI美女」などとうたって1枚数千円で販売され、水谷容疑者は、およそ1年間で1000万円余りを売り上げていたということです。

調べに対し、いずれも行為を認め、このうち水谷容疑者は、「ポスター販売は利益率が高いと知って始めた」などと供述し、菅沼容疑者は、「ポスター販売を事業の1つにしよう」と思い、作成方法は独学で学んだ」などと供述しているということです。

「ディープフェイク」と呼ばれるAIを使った技術で作成した実在する人物や架空の人物の性的な画像がネット上に氾濫しているとして課題となっていますが、警視庁によりますと、生成AIで作成したわいせつ物の販売が摘発されるのは全国で初めてだということです。

AIを使用したわいせつ画像（無修正）を生成し、販売が摘発された初事例

性的ディープフェイクを生成するA Iに関する調査

性的ディープフェイクが生成可能なA Iについて

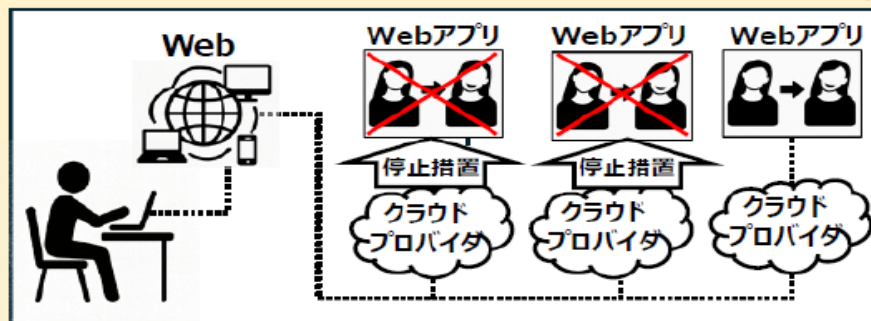
第1回人工知能戦略本部資料

性的ディープフェイクが生成可能な画像・動画生成サービス（アプリ）として、以下を確認。

①Webアプリ （Web上で動作するアプリ）

- 性的ディープフェイクが生成可能とされていた約50種のWebアプリを調査した結果、現在も性的ディープフェイクを生成可能なものは数件のみであった。
- これは、サービスの基盤となるクラウドプロバイダが規約に違反するとして利用停止措置または機能停止勧告したことによる。（図1）
- 現在も生成可能なWebアプリにおいては、実際の写真と見分けがつかない性的画像・動画を容易に生成することができた。

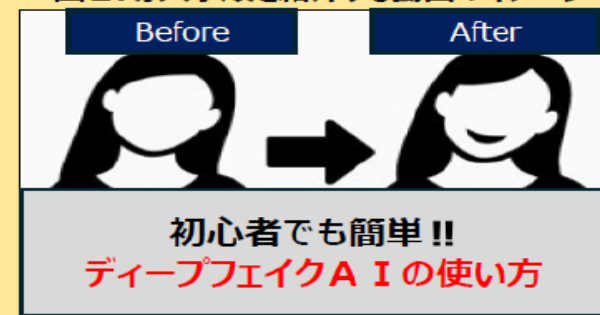
図1. プロバイダによる性的ディープフェイクのWebアプリ利用停止措置



②ダウンロードアプリ （利用者がP Cにダウンロードし利用するアプリ）

- 画像・動画生成可能なダウンロードアプリが少なくとも1万種以上Webで公開されており、性的ディープフェイクが生成可能なものも確認。
- Webアプリと比べて導入の手間が多く、利用時に一定のI Tスキルが必要であるものの、導入・利用マニュアルが出回り（図2）、中学・高校生でも容易に利用可能な状況にあることを確認。
- また、一度ダウンロードされたアプリは継続的に利用可能。

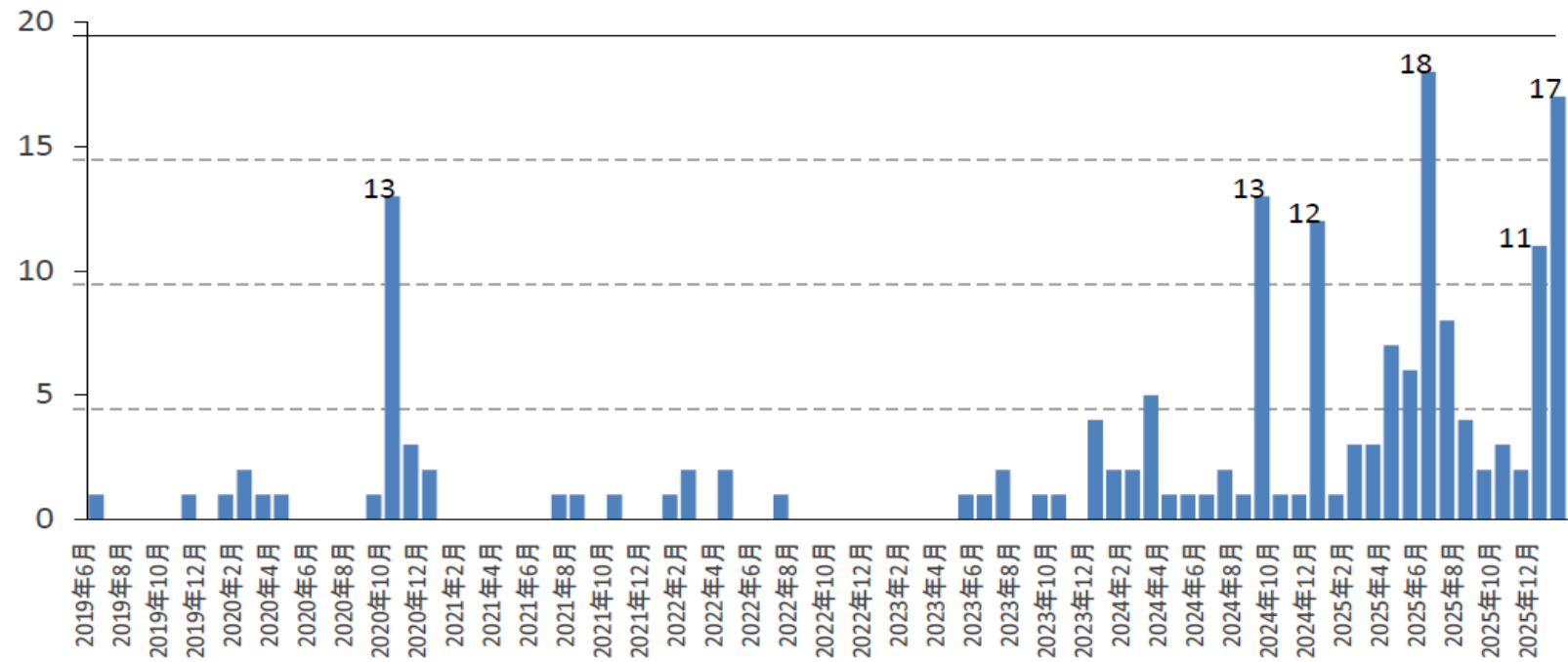
図2. 導入手順を紹介する動画のイメージ



性的ディープフェイクによる被害等の状況

- 性的ディープフェイクに関する報道件数は拡大傾向。
- 被害者が、著名人から一般人に移行。認知件数も増加傾向。
- 主にメンバを限定されたクローズドなWebコミュニティで拡散されている。
- 性的ディープフェイクの拡散状況の全体像を把握することは困難。

ディープフェイクポルノに関する記事数



Kanagawa Pre

- 令和 8 年 2 月 3 日時点。
- 見出し・本文・キーワード・分類語を対象とした右記の検索条件での該当記事172記事。「ディープフェイク AND (ポルノ OR 性的)」
- 検索方式は「すべての語を含む」、一致方式は「完全一致」、同義語展開は「する」、シソーラス展開は「しない」

生成AIポルノ等の被害にあったら

性犯罪・性暴力の被害にあったら

性犯罪・性暴力の被害にあった場合の相談窓口



 かならいん公式

- かながわ性犯罪・性暴力被害者ワンストップ支援センター「かならいん」は、性犯罪・性暴力の被害にあわれたあなたをサポートします。ひとりで悩まず、まずは「かならいん」にご相談ください。（外部委託）（所属：[くらし安全交通課](#)）

まずは「かならいん」へご相談ください

[かながわ性犯罪・性暴力被害者ワンストップ支援センター「かならいん」](#)では、性別・年齢等問わず、性犯罪・性暴力の被害にあわれた方やそのご家族からのご相談をお受けしています。

電話相談 **#8891**（はやくワンストップ）

24時間365日いつでも電話相談を受け付けています。

#8891（はやくワンストップ）は、発信場所から最寄りのワンストップ支援センターにつながる通話料無料の全国共通番号です。（NTTひかり電話からは、0120-8891-77（通話料無料））

全国共通番号は、一部のIP電話等からはつながりません。その際は、**045-322-7379**（通話料がかかります）へお電話ください。

男性及びLGBTs被害者のための専門相談ダイヤル

045-548-5666（毎週火曜日16時から20時 祝休日・年末年始を除く）

// 相談受付時間

毎週 火曜日・木曜日・金曜日・日曜日

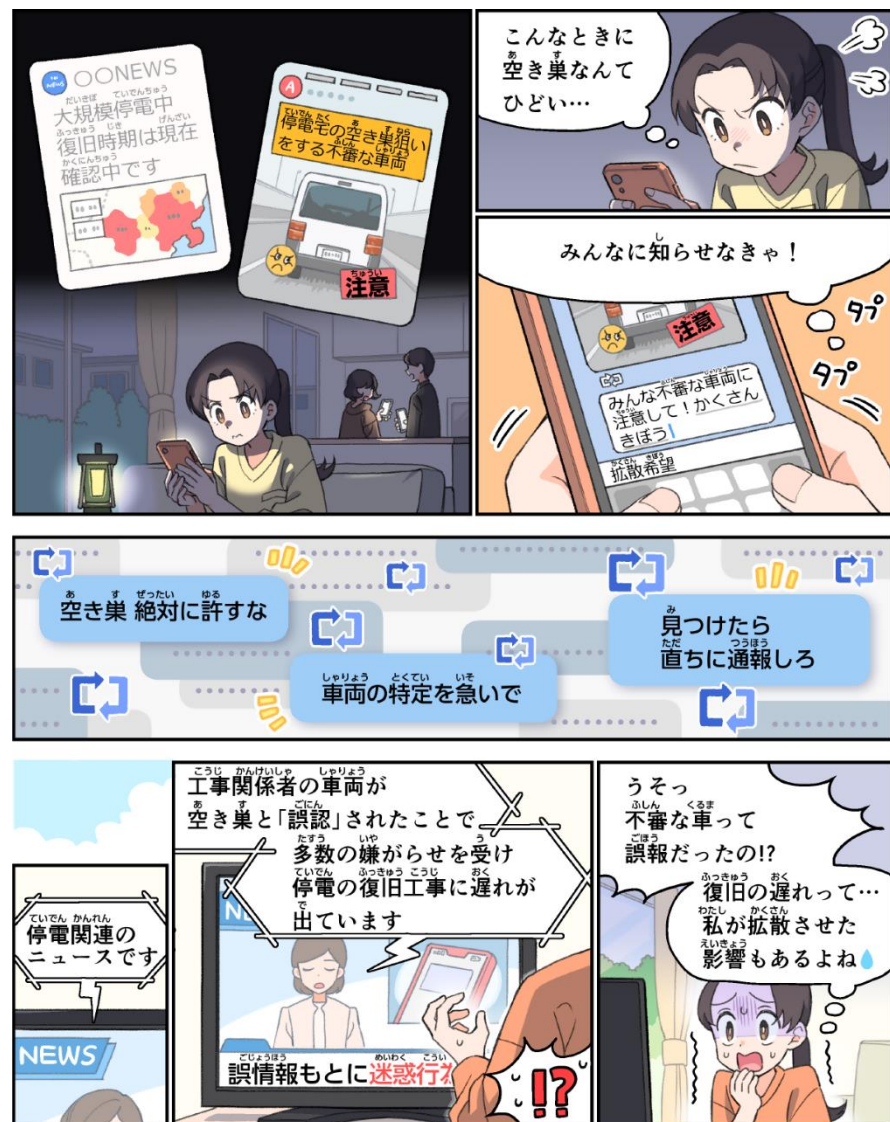
16時～21時



生成AIの課題・リスク

	リスク	事例
従来型AIから存在するリスク	バイアスのある結果及び差別的な結果の出力	● IT企業が自社で開発したAI人材採用システムが女性を差別するという機械学習面の欠陥を持ち合わせていた
	フィルターバブル及びエコーチェンバー現象	● SNS等によるレコメンドを通じた社会の分断が生じている
	多様性の喪失	● 社会全体が同じモデルを、同じ温度感で使った場合、導かれる意見及び回答がLLMによって収束してしまい、多様性が失われる可能性がある
	不適切な個人情報の取扱い	● 透明性を欠く個人情報の利用及び個人情報の政治利用も問題視されている
	生命、身体、財産の侵害	● AIが不適切な判断を下すことで、自動運転車が事故を引き起こし、生命や財産に深刻な損害を与える可能性がある ● トリアージにおいては、AIが順位を決定する際に倫理的なバイアスを持つことで、公平性の喪失等が生じる可能性がある
	データ汚染攻撃	● AIの学習実施時及びサービス運用時には学習データへの不正データ混入、サービス運用時ではアプリケーション自体を狙ったサイバー攻撃等のリスクが存在する
	ブラックボックス化、判断に関する説明の要求	● AIの判断のブラックボックス化に起因する問題も生じている ● AIの判断に関する透明性を求める動きも上がっている
	エネルギー使用量及び環境の負荷	● AIの利用拡大により、計算リソースの需要も拡大しており、結果として、データセンターが増大しエネルギー使用量の増加が懸念されている
生成AIで特に顕在化したリスク	悪用	● AIの詐欺目的での利用も問題視されている
	機密情報の流出	● AIの利用においては、個人情報や機密情報がプロンプトとして入力され、そのAIからの出力等を通じて流出してしまうリスクがある
	ハルシネーション	● 生成AIが事実と異なることをもっともらしく回答する「ハルシネーション」に関してはAI開発者・提供者への訴訟も起きている
	偽情報、誤情報を鵜呑みにすること	● 生成AIが生み出す誤情報を鵜呑みにすることがリスクとなりうる ● ディープフェイクは、各国で悪用例が相次いでいる
	著作権との関係	● 知的財産権の取扱いへの議論が提起されている
	資格等との関係	● 生成AIの活用を通じた業法免許や資格等の侵害リスクも考えうる
	バイアスの再生成	● 生成AIは既存の情報に基づいて回答を作るため既存の情報に含まれる偏見を増幅し、不公平や差別的な出力が継続/拡大する可能性がある

ネット上のフェイクニュース



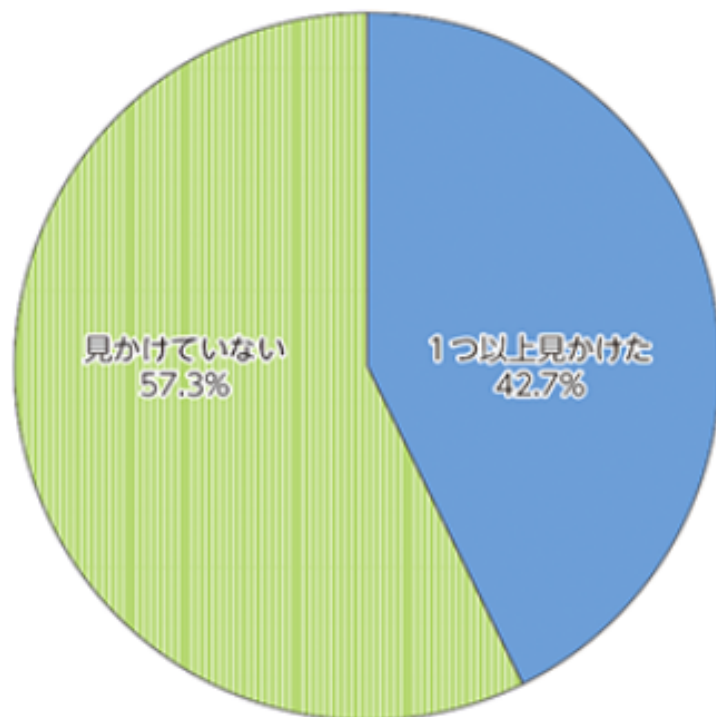
ネット上のフェイクニュース

(3) 真偽不確かな情報との接触

X(旧Twitter)等のSNSは、若年層を中心に震災時の安否確認や情報収集に一定程度寄与した一方、こうしたSNS上では、真偽不確かな情報が拡散し、混乱をもたらした。

こうした真偽不確かな情報について、SNS上で1つ以上「見かけた」と回答した者の割合は42.7%であり、SNSの種類別にみると、X(旧Twitter)の割合が高くなっている。

図表 I-2-1-12 SNS上で真偽不確かな情報を見かけた割合



(出典)総務省(2024)国内外における最新の情報通信技術の研究開発及びデジタル活用の動向に関する調査研究

ネット上のフェイクニュース

次に、見かけた情報の確からしさについてどう感じたかを尋ねてみると、「真偽がわからないと感じた」と回答した割合は各項目で概ね6.5割程度であり、さらにそのうち、「確認しようとした」割合は約3〜4割となっていた。

また、実際に真偽の確認を行った者に確認方法を尋ねたところ、「情報の発信源(組織や人物)を確認した」割合が最も高く(37.6%)、公的な情報、報道機関等による情報を確認した割合はそれぞれ約3割となっていた。

図表 I-2-1-14 真偽不確かな情報をどのように確認したか



こうした情報について1つ以上見かけたと回答した者のうち、1つ以上を「知人へ共有、または不特定多数の人へ拡散したことがある」と回答した者は25.5%であり、その理由としては、「他の人にとって役に立つ情報だと思った」、「その情報が興味深かった」、「その情報が間違っている可能性がある」と注意喚起をしようと思った」といった回答が一定の割合挙げられた。

ネット上のフェイクニュース

「ライオン放たれた」ほかネット偽情報135件 熊本地震で警察記録

飯島健太 2025年1月12日 10時00分



[list](#)



ツイッターに投稿されたライオンの画像＝熊本県警提供 

最大震度7を2回観測した2016年の熊本地震で、熊本県警が把握したインターネット上の「流言飛語」が135件に上っていたことが分かった。朝日新聞の情報公開請求に対し、開示された文書に記録されていた。

熊本地震は16年4月14日に前震、同16日に本震が発生した。文書によると、ネット上では地震後、デマや真偽の不確かな情報が見られた。県警は、被災者の不安や混乱があおられる懸念から「情報対策」に取り組んだという。

その結果、県警は5月8日までに計135件の流言飛語を把握した。主な例に「モールで火災」「動物園からライオンが放たれた」「井戸に毒」「富士山の大噴火が起こる」「迷彩服を着たグループが空き巣」「謎の放射線」などがあつた。

[PR]

このうち、ツイッター（現X）に「ライオンが放たれた」とウソの投稿をして熊本市動物園の業務を妨害したとして、県警は7月、当時20歳の男性を偽計業務妨害の疑いで逮捕した。男性はその後、不起訴処分（起訴猶予）になった。

熊本地震の際に、「ライオンが放たれた」とウソの投稿をして逮捕された事例

Check 1

☑ **情報源** はある？



- その情報は**どこから**、**いつ**発信されたものですか？信用できますか？
- **根拠**となるモノは今も存在していますか？消えていませんか？
- 情報源が「**海外の**」ニュースや論文の場合、あなたはその情報源を**確認、理解**していますか？

Check 2

☑ その分野の **専門家**？



- その情報は、**専門知識**や必要な**資格**を持った人が、責任を持って発信しているものですか？
- その人は過去、二セ・誤情報を発信して**批判**されていませんか？
- その人は関連する情報や商品を**売って**いませんか？

Check 3

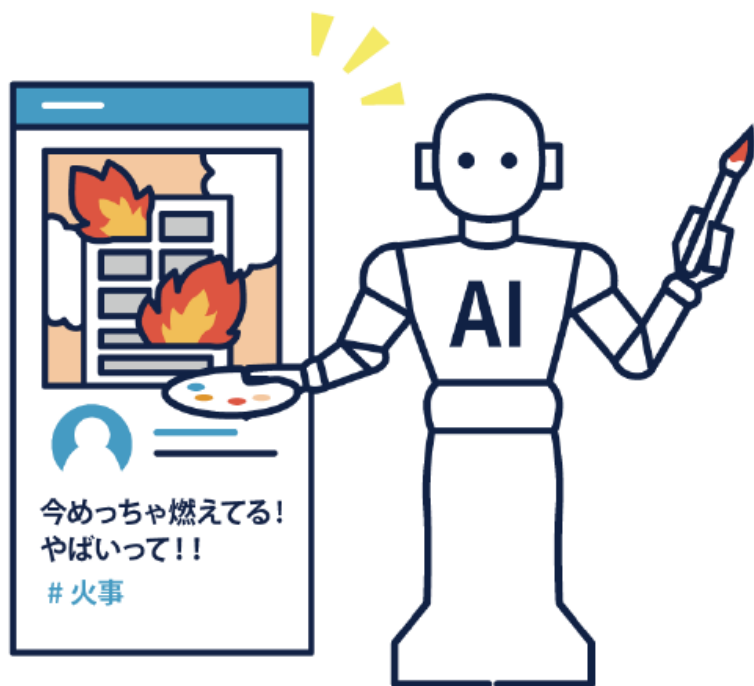
☑ **他では** どう言われている？



- その情報について
他の人や**他のメディア**は
どのように言っていますか？
- その人の意見に
反論している人はいませんか？
- **別の内容**で報じているメディアや、
誤りであることを指摘している
メディアはありませんか？

Check 4

☑ その画像は **本物**？



- 臨場感のある画像が添えられているから？
それだけで「本当」だと判断して大丈夫ですか？
- その画像は**生成AI**によって作成されたものではありませんか？
- その画像は勝手に盗用された、全く**無関係のもの**ではありませんか？

アルゴリズムとは

「自分で選んだ」
と思っているけど...



アルゴリズム



アルゴリズムが選んだ
コンテンツの中から
選んでいる場合も！



その人が好むコンテンツや記事を^{ぶんせき}分析し、どんな^{どうこう}投稿を見せると
楽しんでもらえるか、何を表示するかを決めるアルゴリズムにより、
その人の好みに合ったコンテンツが次々と表示される。



SNSなどのサービス提供者は、サービスの利用時間を増やし、見てもらえる広告を増やして、広告収入が得られるよう様々な工夫をしている。

そのひとつが **アルゴリズム**。

米陪審、幼少期のSNS依存でIT2社に賠償評決 「画期的な判断」 類似訴訟相次ぐ

【ワシントン＝塩原永久】米西部カリフォルニア州ロサンゼルス在地裁の陪審は25日、未成年者のSNS依存を巡る訴訟で、IT大手メタとユーチューブを傘下に持つグーグルの責任を認め、2社に計600万ドル（約9億5千万円）の賠償を命じる評決を下した。米メディアが報じた。

米国ではSNS運営企業を訴える類似の訴訟が相次いでおり、巨大IT企業には打撃になる。メタとグーグルは不服とする声明を発表しており、控訴する公算が大きい。

SNS運営企業は通信品位法230条のもと、投稿される「内容」に関する幅広い免責が認められている。今回の訴訟で原告側は、SNS依存につながるアルゴリズムなどの「仕組み」に焦点をあてて運営企業を追及。賠償責任を認める「画期的な判断」（米紙）となった。

報道によると、原告の女性（20）は、6歳からユーチューブの利用を開始。1日に16時間を費やすなどSNS中毒となり、メンタルヘルスに支障をきたした

陪審団は、SNSの仕組みが未成年に悪影響を及ぼす恐れを把握しながら、運営企業が危険性の周知などの必要な対応を怠ったとした。

評決は賠償額の7割に当たる420万ドルをメタに、残りの180万ドルをグーグルに支払うよう命じた。

裁判では、メタのザッカーバーグ最高経営責任者（CEO）が証人として出廷し、SNSの中毒性を否定した。以前は社内で利用者の滞在時間を長くする目標を設定していたが、取りやめたとともに反論した。

メタを巡っては、米西部ニューメキシコ州の裁判所の陪審も24日、若者をSNS上で性被害などから守る対策を怠ったとし、メタに3億7500万ドルの支払いを命じる評決を下した。



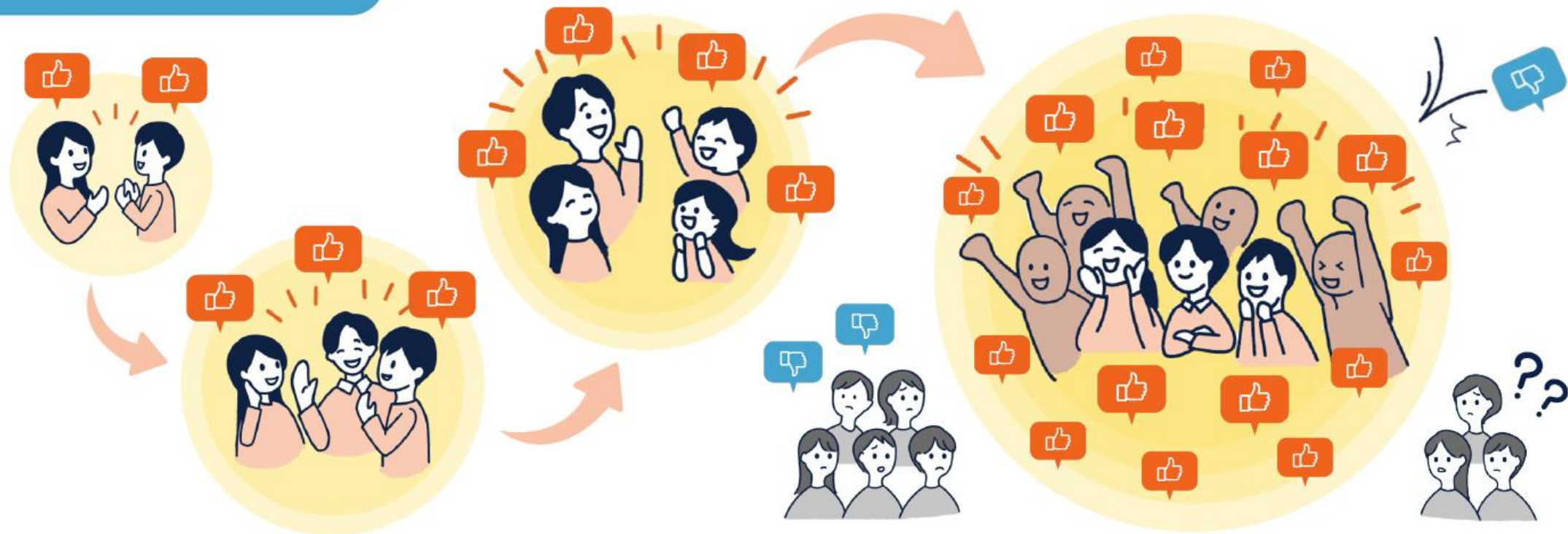
米メタ（旧フェイスブック）のロゴ（ロイター＝共同）

出典：産経新聞ホームページ 2026年3月26日



私たちの多くは好きな食べ物を選び好きな音楽を聴く。
SNSや動画サイトでも、わざわざ興味のないジャンルや、
嫌いなアカウントはたいていフォローしない。

エコーチェンバーとは



私たちは「似た考え方」「同じ趣味^{しゅみ}」の人とつながりやすく、自然と**考え方が**
近いグループが作られ、会話、議論、判断^{たが}が互いに強化され、偏っていく。

これが**エコーチェンバー**。

【エコーチェンバー】

エコーチェンバーとは、同じ意見を持つ人々が集まり、自分たちの意見を強化し合うことで、自分の意見を間違いないものと信じ込み、多様な視点に触れることができなくなってしまう現象を指す。
出典：総務省 令和6年版情報通信白書（p.46脚注1）<https://www.soumu.go.jp/johotsusintokei/whitepaper/ja/r06/pdf/n1410000.pdf>

「ニセ・誤情報」に気付かない人は 「フィルターバブル」に囲まれている可能性も…

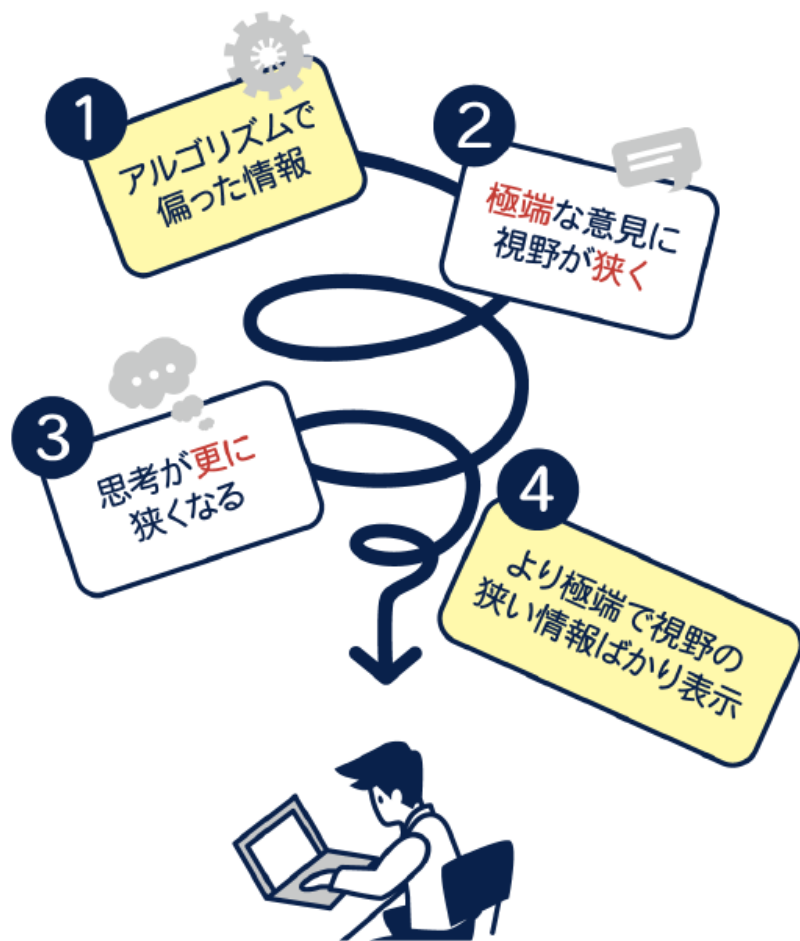


ネットニュース、SNS、検索サービスなどには、**その人が欲しがりそうな情報**を分析し**同じような情報**を表示する「アルゴリズム」と呼ばれる機能があります。

「アルゴリズム」によって知らぬ間にかたよった情報にばかり接する現象を**「フィルターバブル」**と言います。

その結果、それらの情報が**世の中の標準だと誤解**する恐れがあります。

「フィルターバブル」は気付けない！



アルゴリズムは考え方、好みを分析し、その人が心地よいと感じる情報ばかり洪水のように流し込めます。その結果、

「物事を極端にとらえ、狭い視野で考える人」となっていきます。

そして心地よい情報ばかり読むことになっても、フィルターバブルは泡のように透明なので本人はそれに気が付きません。

そんな状況に「ニセ・誤情報」が入り込めば...
何が起きるか分かりますよね？

／ フィルターバブルやエコーチェンバーによって起こりうる対立 ／



他者との違いを認めない、自分と異なる意見は耳に入らないといった
悪影響が、様々な分野で起きていると指摘・懸念されている。

政治的な二セ・誤情報の具体例

根拠のない情報が選挙に影響を与える可能性がある



選挙が近づくと、特定の政治家に対して、

「外国と通じている」

「増税主義者だ」

「もうすぐ逮捕される」

といった根拠のない情報を流し、世間に誤った印象を植付けるといった事例が起きることがあります。

それらが、**一人ひとりの判断に影響**を与えてしまう可能性もあるのです。

出典：総務省 二セ・誤情報にだまされないために

4 世の中があらぬ方向に進んでしまう可能性も…



2020年の米大統領選では

「ウィスコンシン州の投票率が
200%を超えた、バイデン氏による
不正が行われた証拠だ」などの

ニセ・誤情報拡散 がきっかけとなり、
のちにトランプ氏支持者が

ワシントンの連邦議会を襲う

という前代未聞の事件につながったと
指摘されています。



2021年1月6日、「ストップ・ザ・スチール」と書かれたプラカードなどを手に、ワシントンに集まったトランプ大統領（当時）の支持者=AP 

2021年1月6日、トランプ米大統領（当時）の支持者らがワシントンの連邦議会議事堂を襲撃した事件は、SNSで広まった誤情報が、現実の危険を招いた例だった。事件は、フェイスブック（FB、現メタ）がトランプ氏のアカウントを停止するきっかけともなった。しかし、朝日新聞が入手した同社の内部文書からは、事件前からトランプ氏の支持者らの動きに危機感を抱きながら、対応が後手に回った様子が浮かんた。

FBにとって、20年11月3日の米大統領選への対応は課題だった。前回16年の大統領選ではFBを通じて虚偽情報が拡散されたと指摘され、20年の大統領選でもトランプ氏の発言などに何度も警告などを出した。効果あってか、投票日は大きな混乱がなかった。

だが、開票が進みトランプ氏の劣勢が判明すると、根拠のない「不正投票」の主張がすぐ拡散した。選挙直後に支持者が始めた「ス

トップ・ザ・スチール（盗みを止めろ）」という運動を検証したFBの内部文書は「大きな問題がなく選挙が過ぎたという満足感、怒りの暴言と陰謀論の増加で次第に変化した」と振り返っている。

ニセ・誤情報拡散が
きっかけとなり、
ワシントン連邦議会を襲う
事件につながった事例

- 1 SNSや動画サイトなどから「受け取る」情報だけに頼らず、
自分から情報を探し出す（＝「取りに行く」）力をつける
- 2 SNSや動画サイトなどでは、人々が受け取る情報が
操られやすい環境かんきょうにあることを意識する
- 3 SNSや動画サイトなどから得られる情報だけで決めつけず、
自分が見ている情報が世の中の全てではないことを意識する

生成AI活用に当たって注意すべきポイントは？

情報の正確性

- ✓ 無意識のうちに合理的ではない行動、偏った判断をすることがあるという意識を持つ
- ✓ チェックリストを用いて真偽を判断する
- ✓ 安易に拡散しない / 拡散したいときはひと呼吸おく

情報流出

- ✓ 生成AIサービスの規約を確認する（商用利用可否、損害発生時の責任所在等）
- ✓ 個人情報や機密情報の入力が必要最小限にする
- ✓ 生成AIに入力したデータを学習に使わせないように設定する

知的財産権の侵害

- ✓ 既存のものや実在の人物に似たものを生成するような指示入力を避ける
- ✓ 生成物が既存のものや実在の人物に類似している場合、利用をやめる / 権利者から許諾を取得後に利用する / 既存のものと類似しないよう大幅に加工する

活用者としてのモラル

- ✓ 本来自分が行うべきことまで生成AI任せにしない
- ✓ 生成AIが作った偏見のある回答を使用しない
- ✓ 生成AIを非倫理的な行為や犯罪に悪用しない