

第6章 米国

1. 政府生成AI関連規制に係る当局の位置づけ及びその概要

(1). 生成AI規制に係る当局の概要

米国で消費者保護に関する生成 AI の規制を主導するのは、米国連邦取引委員会（Federal Trade Commission：FTC）²⁴⁵である。FTCは1914年の連邦取引委員会法によって設立された独立機関で、不公正又はぎまんの取引慣行を防止する責務を負っている。設立当初は独占禁止（競争保護）を目的としたが、その任務は徐々に拡大され、企業による不公正な競争方法や消費者に対するぎまんの・不公正な行為を取り締まることが主要な役割となった。

FTCの活動目的は、企業の違法又は有害な行為から国民を保護し、健全な市場環境を維持することにある。具体的には、不公正・ぎまんの商慣行の排除や、消費者が十分な情報に基づいて選択できる市場の促進などを目的としている。FTCは法執行機関としての権限を持ち、企業や個人の違法行為を調査し、法的措置（差止命令や制裁措置の追求など）を行う。また消費者や事業者への教育・啓発、市場動向の調査、他機関への政策提言といった機能も担い、広範な手段で消費者保護に取り組んでいる。昨今では生成 AI を含む先端技術による新たなリスクにも注目しており、AI が関与する詐欺や虚偽表示などにも既存の消費者保護法を積極的に適用する方針を明確にしている²⁴⁶。

(2). 当局の位置づけ

FTCの権限は、主に連邦取引委員会法第5条²⁴⁷に定められており、これにより「不公正又はぎまんの取引慣行」が厳しく禁止される。さらに、同法は競争政策に関する規定も含むため、反トラスト法との連携によって市場全体の健全性が担保される仕組みとなっている。こうした法的根拠は、生成 AI が引き起こす消費者被害についても既存の法律の枠内で対応可能であるとの見解を裏付けるものである²⁴⁸。

FTCは、5名の委員（Commissioners）によって構成され、そのうち1名が委員長（Chair）を務める。委員は大統領により指名され、上院の承認を経て任命されるため、政治的な中立性が求められる²⁴⁹。

意思決定は、委員会全体の多数決方式により行われ、各部門が連携して迅速かつ効果的な対応を行う体制が整えられている。こうした組織構成は、生成 AI による新たな消費者被害にも柔軟に対応するための重要な基盤となっている。

²⁴⁵ <https://www.ftc.gov/>

²⁴⁶ <https://www.ftc.gov/news-events/news/press-releases/2023/04/ftc-chair-khan-officials-doj-cfpb-eeoc-release-joint-statement-ai>

²⁴⁷ <https://uscode.house.gov/view.xhtml?req=granuleid%3AUSC-prelim-title15-chapter2-subchapter1&edition=prelim>

²⁴⁸ <https://www.ftc.gov/policy/advocacy-research/tech-at-ftc/2025/01/ai-risk-consumer-harm>

²⁴⁹ <https://www.jftc.go.jp/kokusai/worldcom/alphabetic/u/america.html>

(3). 関連省庁・機関の概要

① 消費者金融保護局 (Consumer Financial Protection Bureau : CFPB) ²⁵⁰

CFPB は、2010 年のドッド＝フランク法に基づいて設置され、金融サービス分野における消費者保護を専門に担当する。金融機関の不正な取引慣行や不当な貸付審査、クレジットカード会社による不正行為などに対して厳格な監督を実施する。近年、生成 AI を利用した融資審査やチャットボット対応の自動化に伴うリスクに対しても、既存の規制枠組みが適用されるべきとの見解が示され、消費者の権利保護を強化する方向性が打ち出されている²⁵¹。

② 商務省関連機関 (NTIA及びNIST)

商務省傘下の NTIA は、インターネット及び通信分野の政策立案を担い、AI 技術の説明責任や透明性向上に向けた取組を推進している²⁵²。一方、NIST は業界標準や技術ガイドラインの策定を通じ、生成 AI の信頼性・安全性・公正性を評価する枠組みを提供している。2023 年には NIST が AI リスク管理フレームワークを発表し、企業が AI システムの安全性を確保するためのベストプラクティスが示された。これにより、消費者保護の視点からも生成 AI の技術的信頼性が担保される仕組みが整備されつつある。

③ 連邦通信委員会 (FCC)

FCC は、通信分野全般の規制を担う独立機関であるが、近年、生成 AI が関与する詐欺的通信やディープフェイクを利用した詐欺行為への対策にも注目している。2024 年には、消費者保護を目的とした FCC の消費者諮問委員会が再編成され、AI 技術がもたらすリスクとその対策に関する議論が活発に行われる体制が整備された²⁵³。通信インフラとサービスの安全性確保を通じ、消費者のプライバシーや安全が守られるよう、FTC との連携が図られている。

④ 国家人工知能諮問委員会 (National Artificial Intelligence Advisory Committee : NAIAC) ²⁵⁴

NAIAC は、2020 年国家人工知能イニシアチブ法 (National Artificial Intelligence Initiative Act of 2020) ²⁵⁵に基づき、商務長官の指名を受けた産学官の専門家によって構成される高水準な諮問機関である。大統領及びホワイトハウス内の AI 政策調整オフィスに対して、AI の倫理、透明性、説明責任などを含む総合的な提言を行い、生成 AI がもたらす消費者被害に対するリスク管理の枠組みの策

²⁵⁰ <https://www.consumerfinance.gov/>

²⁵¹ <https://www.consumerfinance.gov/about-us/newsroom/cfpb-federal-partners-confirm-automated-systems-advanced-technology-not-an-excuse-for-lawbreaking-behavior/>

²⁵² <https://www.sidley.com/en/insights/newsupdates/2023/11/president-biden-signs-sweeping-artificial-intelligence-executive-order>

²⁵³ <https://docs.fcc.gov/public/attachments/DOC-400597A1.pdf>

²⁵⁴ <https://www.nist.gov/itl/national-artificial-intelligence-advisory-committee-naiac>

²⁵⁵ <https://uscode.house.gov/view.xhtml?path=/prelim@title15/chapter119&edition=prelim>

定に寄与している。NAIACの提言は、各規制当局の施策やガイドラインの見直しに反映され、政府全体でのAIガバナンスの強化が図られている。

⑤ FCC消費者諮問委員会（Consumer Advisory Committee：CAC）²⁵⁶

FCCの消費者諮問委員会は、通信分野におけるAI技術の活用とそれに起因する消費者リスクに関して、業界団体や消費者団体、有識者から構成される助言機関である²⁵⁷。これにより、生成AIを用いた詐欺行為やディープフェイクによる不正利用への対応策が検討され、実務的なガイドラインの策定に寄与している。

²⁵⁶ <https://www.fcc.gov/consumer-advisory-committee>

²⁵⁷ <https://docs.fcc.gov/public/attachments/DOC-400597A1.pdf>

2. 消費者保護に関する生成 AI に関連する法規及びその所管状況や基本計画等

(1). 関連法令

① 連邦取引委員会法 (Federal Trade Commission Act) ²⁵⁸

連邦取引委員会法は、1914 年 9 月 26 日に制定され、その後 1938 年の Wheeler - Lea 改正法により消費者保護規定が強化された法律である²⁵⁹。同法は、連邦取引委員会 (FTC) が専管する独立行政機関としての権限の根拠となっており、同法第 5 条において「不公正又はぎまんの取引慣行」を禁止する規定を有する。連邦取引委員会法の目的は、消費者を欺く行為や不公正なビジネス慣行から国民の利益を保護し、公正かつ自由な市場環境を確保することにある。

同法は、州際通商に関与する企業や業種全般に適用され、生成 AI を用いたフェイク広告、詐欺的な情報提供、さらには AI によるディープフェイクや不正ななりすまし行為といった事例にも、既存の法的枠組みを適用して取り締まる機能を持つ。生成 AI 規制に関しては、「AI だからといって既存法の適用から免れることはない」という原則を強調しており、FTC は 2023 年以降、AI 技術を悪用した詐欺行為に対して積極的な法執行措置を講じる方針を示している²⁶⁰。これにより、生成 AI を利用した消費者ぎまんやフェイクニュース拡散は、従来の不公正取引慣行として法的に規制されることとなり、企業に対する差止命令や民事罰の請求などが実施される仕組みとなっている。

② 2010年7月21日成立のドッド＝フランク金融規制改革法の第10章「消費者金融保護法」 (Consumer Financial Protection Act) ²⁶¹

消費者金融保護法は、2010 年 7 月 21 日に成立したドッド＝フランク金融規制改革法の一部として制定された法律であり、同法に基づいて設置された消費者金融保護局 (CFPB) が専管する。当該法の目的は、金融商品やサービスにおける不公正、ぎまんの、又は濫用的な取引慣行 (Unfair, Deceptive, or Abusive Acts or Practices : UDAAP) の禁止を通じ、消費者の金融的利益を保護することである。

消費者金融保護法は、銀行やノンバンクを含む金融サービス提供者全般に適用され、生成 AI が利用される場合においても、例えば AI を用いた融資審査や信用評価におけるバイアス、誤情報提供による消費者被害に対して厳格な監視と規制を実施する。当該法においては、企業に対して事前の影響評価やアルゴリズムの透明性の確保、さらには不正な審査方法の是正措置が義務付けられており、AI

²⁵⁸ 条文 : <https://uscode.house.gov/view.xhtml?req=granuleid%3AUSC-prelim-title15-chapter2-subchapter1&edition=prelim>

²⁵⁹ 同上

²⁶⁰ <https://www.ftc.gov/news-events/news/press-releases/2024/09/ftc-announces-crackdown-deceptive-ai-claims-schemes>

²⁶¹ 『DODD-FRANK WALL STREET REFORM AND CONSUMER PROTECTION ACT』 581P-739P : <https://www.congress.gov/111/plaws/publ203/PLAW-111publ203.pdf>

による与信審査の結果が不当な差別や偏りを生じた場合には、平等信用機会法などと連携して厳正な対応が図られる。CFPB は、AI 技術がもたらす金融サービス分野のリスクを、従来の規制枠組みの中で包括的に管理し、金融市場における消費者保護の水準を向上させるための重要な役割を担っている²⁶²。

③ 児童オンラインプライバシー保護法（Children's Online Privacy Protection Act : COPPA）²⁶³

児童オンラインプライバシー保護法は、1998 年 10 月に制定され、2000 年 4 月に施行された法律であり、連邦取引委員会（FTC）が専管する²⁶⁴。当該法は、13 歳未満の児童の個人情報をオンライン上で保護することを目的としており、ウェブサイト運営者やオンラインサービス提供者に対して、児童から個人情報を収集する際の保護措置（親の同意取得、プライバシーポリシーの明示など）を義務付けるものである。

生成 AI に関しては、例えば対話型 AI やチャットボットが児童向けサービスとして提供される場合、又は児童が利用可能なプラットフォーム上で生成 AI が用いられる場合に、当該法の適用対象となる。児童オンラインプライバシー保護法は、児童のプライバシー権の確保と安全なオンライン環境の提供を通じて、将来的に生成 AI が引き起こす潜在的なプライバシー侵害や個人情報の不正利用といった問題に対しても、厳格な規制の枠組みを提供するものである²⁶⁵。

（2）. 基本計画の概要

2021 年以降、米国政府及び関連機関は、生成 AI 技術の急速な普及に伴い、従来の法規制の枠組みを越えた新たな規制・政策戦略を検討し始めた。まず、ホワイトハウス科学技術政策局（Office of Science and Technology Policy : OSTP）²⁶⁶が発表した「Blueprint for an AI Bill of Rights（AI 権利章典の青写真）²⁶⁷」は、AI システムが守るべき基本的な原則として、個人のプライバシー保護、説明責任、公平性、有害な偏見からの保護、人的介入の機会の確保を掲げたものである。この指針は法的拘束力を有さないものの、各連邦機関や企業に対して、生成 AI を含む自動化システムの安全かつ倫理的な利用を促すためのガイドラインとして機能し、今後の規制の土台となることが期待される。

また、国立標準技術研究所（NIST）が 2023 年初頭に発表した「Artificial Intelligence Risk Management Framework（AI RMF 1.0）²⁶⁸」は、生成 AI を含むあらゆる AI システムのリスク評価及び管理のための技術的基盤を提供するものである。NIST は、このフレームワークを通じて、企業

²⁶² <https://www.consumerfinance.gov/about-us/newsroom/cfpb-comment-on-request-for-information-on-uses-opportunities-and-risks-of-artificial-intelligence-in-the-financial-services-sector/>

²⁶³ 条文： <https://uscode.house.gov/view.xhtml?req=granuleid%3AUSC-prelim-title15-section6501&edition=prelim>

²⁶⁴ 同上

²⁶⁵ <https://www.ftc.gov/policy/advocacy-research/tech-at-ftc/2025/01/ai-risk-consumer-harm>

²⁶⁶ <https://www.usa.gov/agencies/office-of-science-and-technology-policy>

²⁶⁷ <https://bidenwhitehouse.archives.gov/ostp/ai-bill-of-rights/>

²⁶⁸ Artificial Intelligence Risk Management Framework

や組織が自発的にリスク管理を実施できるようにするだけでなく、連邦政府が将来的に規制基準として採用する可能性を視野に入れている。これにより、生成 AI の透明性や説明責任、さらにはアルゴリズムに内在するバイアスの低減が進むとともに、消費者被害を未然に防ぐ仕組みが整備される見込みである。

さらに、バイデン政権は 2023 年 10 月に包括的な大統領令を発出し、生成 AI を含む高度な AI モデルに対して安全基準の策定やレッドチームによる脆弱性評価、さらには生成コンテンツの識別技術の開発を指示した²⁶⁹。この大統領令は、従来の消費者保護法規が技術革新に対応するための補完的施策として位置付けられており、生成 AI がもたらす詐欺、差別、プライバシー侵害などに対して、各連邦機関が一層厳格な監視と取り締まりを実施するための枠組みを構築するものである。特に、生成 AI によるディープフェイクの作成や不正利用が消費者に与える影響に鑑み、連邦取引委員会（FTC）及び消費者金融保護局（CFPB）は、既存の法規の枠組みを活用しつつ、新たなガイドラインの策定やルール改正に着手している。

さらに、連邦政府の各機関は、AI 技術に関する自主的な企業コミットメントの確保にも注力している。大手 AI 企業に対しては、製品公開前の安全検証、生成コンテンツの識別技術（透かしなど）の導入、さらにはアルゴリズムの説明可能性の向上などを自主的に実施するよう求める動きが見られる。ホワイトハウスはこれらの取組を評価し、今後の規制策定に反映させる意向を示している²⁷⁰。

以上のように、米国における消費者保護関連の生成 AI 規制は、既存の法規の厳格な適用と新たな政策戦略の検討という二層のアプローチを採用している。現時点では、連邦取引委員会法、消費者金融保護法、児童オンラインプライバシー保護法などが生成 AI に由来する不公正・ぎまん行為に対して適用され、さらに議会や大統領令を通じた新たな規制措置の整備が進行中である。今後、AI 技術の進展とともに、生成 AI に関する具体的な規制の明確化及び実効性のある法制度の整備が、消費者保護の観点からますます重要となる。

²⁶⁹ <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>

²⁷⁰ <https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2023/09/12/fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-eight-additional-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/>

3. 消費者保護に関する生成AI関連規制を行う当局が政策立案の基礎としている事実・社会状況等と、それらを把握するための仕組み

(1). 政策立案の基礎とする事実・社会状況

① 生成AIなりすまし行為の増加

連邦取引委員会（FTC）は、生成 AI を活用した詐欺行為が、従来の詐欺手法に比べて迅速かつ大規模に拡散しているとの実態を把握している。例えば、AI を利用しディープフェイクなどによる詐欺について、消費者からの苦情件数が過去数年間で大幅に増加した²⁷¹。FTC が運用する消費者苦情データベース（Consumer Sentinel Network）²⁷²においては、これら詐欺行為に関する報告件数が 2022 年度から 2023 年度にかけて顕著に上昇しており、被害総額が数十億ドルに達していることが確認されている。これにより、消費者保護の観点から生成 AI を悪用した不公正な取引慣行を厳格に取り締まる必要性が浮上した。

② アルゴリズムの偏見・差別に関する問題

連邦政府及び各州政府は、生成 AI が活用される各分野において、アルゴリズムの偏見や差別が実際に問題となっていることをデータで把握している。金融、雇用、住宅、医療などの分野において、AI を用いた審査や判断が、特定の人種や性別に不利益をもたらす事例が確認されている²⁷³。具体的には、信用審査において AI が過去の不平等なデータに基づいて判断する結果、特定のグループに対して不公平な与信条件が課されるケースや、採用プロセスにおいて AI による選考が不当に差別的な結果を招いた事例が存在する。こうした状況は、従来の差別問題と同様に、生成 AI 技術の利用においても法の厳格な適用が必要であるとの認識の根拠となっている²⁷⁴。

③ プライバシー侵害及び個人データの悪用

生成 AI の学習には大量のデータが必要であり、これにより消費者の個人情報が無断で利用されるリスクが増大している。連邦取引委員会（FTC）は、消費者の音声、映像、テキストなどが生成 AI の訓練データとして用いられ、本人の同意なく拡散されるケースが報告されていることを確認している。特に、音声アシスタントやスマートホーム機器に関連しては、ユーザーの発言が記録・解析される事例があり、これらのデータが第三者に提供された場合、プライバシー侵害及び個人情報の漏洩が発生

²⁷¹ <https://www.ftc.gov/news-events/news/press-releases/2024/02/ftc-proposes-new-protections-combat-ai-impersonation-individuals>

²⁷² <https://www.ftc.gov/enforcement/consumer-sentinel-network>

²⁷³ <https://www.consumerfinance.gov/about-us/newsroom/cfpb-outlines-options-to-prevent-algorithmic-bias-in-home-valuations/>

²⁷⁴ <https://www.consumerfinance.gov/about-us/newsroom/cfpb-federal-partners-confirm-automated-systems-advanced-technology-not-an-excuse-for-lawbreaking-behavior/>

するとの実態が存在する²⁷⁵。

④ 市場規模及び技術動向に関する統計データ

連邦政府は、生成 AI 市場の急速な拡大及びそれに伴う消費者リスクについて、各省庁が独自の調査を実施し、統計データとして収集している。連邦取引委員会（FTC）は、消費者から寄せられる苦情件数や被害額のデータを集計し、生成 AI を悪用した詐欺やフェイク広告の事例が、近年急激に増加していることを示している²⁷⁶。消費者金融保護局（CFPB）もまた、金融分野における AI ツールの利用状況や、その結果としての与信審査のバイアスに関するデータを公表しており、米国の主要銀行におけるチャットボット利用率や、AI システム導入後の消費者苦情の増減が明らかにされている²⁷⁷。これらの統計は、当局が消費者保護政策の必要性を裏付ける客観的な根拠として利用されている。

（２）. 事実・社会状況を把握するための仕組み

① 連邦政府機関によるデータ収集及び監視システム

連邦取引委員会（FTC）は、Consumer Sentinel Network²⁷⁸と呼ばれる苦情データベースを運用し、全国から寄せられる消費者苦情を一元的に収集している。当該ネットワークは、詐欺、フェイク広告、なりすまし、ディープフェイクなどに関する苦情を集計し、統計的な傾向を把握するために活用されている。これにより、生成 AI を悪用した不正行為の件数や被害額の増加がリアルタイムに監視され、必要に応じた法執行措置やガイドライン改正の基礎資料となっている²⁷⁹。また、消費者金融保護局（CFPB）は、金融サービスに関連する苦情及び市場調査を実施している。例えば、消費者金融におけるチャットボットについての調査を実施し、レポートとして報告を行った²⁸⁰。

② 政府による調査研究プロジェクト及び公的レポート

連邦政府は、生成 AI の社会的影響を定量的に把握するため、各省庁が独自の調査研究プロジェクトを実施している。例えば、商務省電気通信情報局（NTIA）は、AI 説明責任ポリシーの策定にあた

²⁷⁵ <https://www.ftc.gov/news-events/news/press-releases/2024/07/ftc-submits-comment-fcc-work-protect-consumers-potential-harmful-effects-ai>

²⁷⁶ <https://www.ftc.gov/news-events/news/press-releases/2024/02/ftc-proposes-new-protections-combat-ai-impersonation-individuals>

²⁷⁷ <https://www.consumerfinance.gov/data-research/research-reports/chatbots-in-consumer-finance/chatbots-in-consumer-finance/>

²⁷⁸ <https://www.ftc.gov/enforcement/consumer-sentinel-network>

²⁷⁹ <https://www.ftc.gov/news-events/news/press-releases/2024/02/ftc-proposes-new-protections-combat-ai-impersonation-individuals>

²⁸⁰ <https://www.consumerfinance.gov/data-research/research-reports/chatbots-in-consumer-finance/chatbots-in-consumer-finance/>

り、意見募集（Request for Comments：RFC）を実施し、1447 件のコメントを受け取った²⁸¹。また、ホワイトハウス科学技術政策局（OSTP）は、2022 年に「AI 権利章典の青写真（Blueprint for an AI Bill of Rights）」を策定するにあたり、一般市民、技術者、学者、有識者から広範な意見を収集した。このプロセスにおいては、オンラインフォームによる意見募集や、パブリックヒアリングが実施され、集められたデータが文書化され、政府全体の AI 政策の基礎資料として採用されている²⁸²。

③ 有識者会議及び諮問委員会の開催

連邦レベルでは、国家人工知能イニシアチブ法に基づいて設置された国家 AI 諮問委員会（NAIAC）が、産学官の専門家を交えた会議を定期的に開催している²⁸³。NAIAC は、生成 AI を含む AI 技術の倫理的利用、透明性、説明責任に関する最新の研究成果及び社会的影響について議論し、その結果を大統領や連邦各省庁に対する提言としてまとめている。これらの会議は、議事録や報告書として公表され、法案や大統領令の立案に反映される。州レベルにおいても、カリフォルニア州政府が州内トップ大学（UC バークレーやスタンフォード大学）と公式提携し、生成 AI の社会影響を共同研究する枠組みを構築している²⁸⁴。

④ 州政府・自治体レベルにおける独自の情報収集及び評価システム

連邦政府のみならず、各州政府や自治体も独自の仕組みを構築して、生成 AI の利用状況及びその影響に関する情報を収集している。カリフォルニア州は、州内企業及び行政機関と連携して州全体で AI 技術の開発、使用、リスクを調査し、州政府内で AI を評価及び導入するための慎重かつ責任あるプロセスを開発するための行政命令を行った²⁸⁵。また、ニューヨーク州では、求人に AI を用いる企業に年次監査報告を義務付け、その報告データを市行政が収集・公開する仕組みが 2023 年より施行された²⁸⁶。これらの取組は、州ごとの実態把握に基づく迅速な対応策の策定を可能にし、連邦政府の規制政策と連動して消費者保護の強化に寄与している。

4. 消費者保護に関する生成AIの最近の動き

（1）. 民間企業における生成AI活用の取組

① Amazon：AIによるレビューの信頼性確保

Amazon は、EC 市場における信頼性確保のため、AI を用いて不正なレビューの検出と要約を行っている。Amazon では、消費者が製品を購入する際に、顧客レビューが大きな役割を果たす。しかし、

²⁸¹ <https://www.ntia.gov/press-release/2023/ntia-receives-more-1400-comments-ai-accountability-policy>

²⁸² [tps://bidenwhitehouse.archives.gov/ostp/ai-bill-of-rights/](https://bidenwhitehouse.archives.gov/ostp/ai-bill-of-rights/)

²⁸³ <https://www.nist.gov/itl/national-artificial-intelligence-advisory-committee-naiac>

²⁸⁴ <https://www.gov.ca.gov/2023/09/06/governor-newsom-signs-executive-order-to-prepare-california-for-the-progress-of-artificial-intelligence/>

²⁸⁵ 同上

²⁸⁶ 『New York City Department of Consumer and Worker Protection』： <https://rules.cityofnewyork.us/wp-content/uploads/2023/04/DCWP-NOA-for-Use-of-Automated-Employment-Decisionmaking-Tools-2.pdf>

近年ではフェイクレビューが増加しており、消費者が正しい判断を下すことが困難になっている。これを防ぐため、Amazon は AI を活用し、不審なレビューのパターンを検出し、該当レビューを削除する仕組みを導入した。2023 年の時点で、Amazon は世界全体で 2 億 5 千万件以上の不審なレビューを削除しており、レビューの信頼性を向上させている²⁸⁷。また、消費者の利便性向上のため、生成 AI を活用し、製品レビューの要点を自動要約する機能も提供している。これにより、消費者は長文のレビューを読むことなく、製品のメリット・デメリットを短時間で把握できるようになった²⁸⁸。

② Microsoft：生成AIの悪用防止対策

Microsoft は、自社の生成 AI が詐欺やフェイクニュースの拡散に悪用されないよう、技術的対策を強化するとともに、法的措置も講じている。2025 年 1 月、Microsoft は、生成 AI のセキュリティ機能を無効化し、不正なコンテンツ生成を可能にするツールを販売していた犯罪グループに対し、米連邦裁判所に提訴を行った。このグループは、AI サービスへの不正アクセスを行い、ポルノや詐欺広告に使用できるディープフェイクを生成するソフトウェアを販売していたとされる Microsoft は、これらの違法行為を防ぐために、該当するアカウントのアクセスを無効化し、裁判所の命令を得て、犯罪に使用されたウェブサイトを閉鎖した²⁸⁹。また、2023 年には生成 AI を悪用したコンテンツの拡散を防ぐための透明性ガイドラインを発表し、業界全体でのルール作りを進めている²⁹⁰。

③ Visa：決済詐欺の検知への生成AI活用

クレジットカード大手の Visa は、オンライン決済における詐欺対策として生成 AI を導入している。特に「列挙型攻撃」と呼ばれる不正アクセス手法では、自動スクリプトを用いてカード番号や有効期限を試行錯誤し、盗難カードを悪用する手口が増えている。Visa は、2024 年 5 月に発表した「Visa アカウント攻撃インテリジェンス (VAAI)」を活用し、リアルタイムで各取引のリスクをスコアリングすることで、不審な取引を即座に検知・ブロックする仕組みを導入した。これにより、誤検知率を 85%削減しつつ、不正取引の検出精度を向上させた。この AI モデルは 150 億件以上の決済データを学習しており、各カード発行会社にも提供されることで、消費者の決済被害を大幅に軽減することが期待されている²⁹¹。

(2). 公共機関における生成AI活用の取組

① デンバー市：市民向け多言語AIチャットボット「Sunny」

デンバー市は、住民向けの行政サービスを強化するため、生成 AI を活用したチャットボット

²⁸⁷ <https://www.aboutamazon.eu/news/customer-trust/how-amazon-is-using-ai-to-detect-fake-product-reviews-and-ensure-authentic-customer-feedback>

²⁸⁸ <https://www.aboutamazon.com/news/amazon-ai/amazon-improves-customer-reviews-with-generative-ai>

²⁸⁹ <https://blogs.microsoft.com/on-the-issues/2025/01/10/taking-legal-action-to-protect-the-public-from-abusive-ai-generated-content/>

²⁹⁰ 同上

²⁹¹ <https://usa.visa.com/about-visa/newsroom/press-releases.releaseId.20661.html>

「Sunny」を導入した。Sunny は、市民からの問合せに 24 時間対応し、英語・スペイン語・ドイツ語・フランス語を含む 72 言語で利用可能である。住民は、ゴミ収集日や違法駐車の情報方法など、市政に関する情報を簡単に取得できる。この取組により、住民の利便性が向上するとともに、市役所の問合せ業務が効率化された²⁹²。

② サンフランシスコ市：311苦情対応へのAI活用

サンフランシスコ市は、市民からの苦情・要望受付システム「311」に生成 AI を活用し、問合せの処理時間を短縮した。AI が市民の投稿内容や画像を解析し、自動で担当部門へ振り分けることで、従来の人手による振り分け作業を大幅に効率化した²⁹³。また、市職員が生成 AI を適切に利用するための「生成 AI 活用ガイドライン」も策定し、個人情報の入力禁止や AI 生成コンテンツの明示など、透明性を確保する取組を進めている²⁹⁴。

③ 音声クローン詐欺の増加

米連邦取引委員会（FTC）は、生成 AI を用いた音声クローン詐欺（ボイスフィッシング）が急増していると警告している。この詐欺は、わずか数秒間の音声サンプルを基に、ターゲットの家族の声を再現し、緊急事態を装って金銭を要求する手口を用いる²⁹⁵。被害者が実際の家族の声と誤認して送金してしまうケースが多数報告されており、FTC や FBI は「本人確認なしに送金しないように」と呼びかけている²⁹⁶。

④ フェイクレビュー・偽プロフィールの氾濫

AI で生成された偽のプロフィール写真や口コミレビューが、EC サイトや SNS で急増している。FTC の調査によれば、企業が AI を利用してやらせレビューを作成するケースもあり、消費者が正確な情報を得る妨げとなっている²⁹⁷。2024 年には、Meta が SNS 上の AI 生成コンテンツにラベルを付与する方針を発表し、透明性の確保を目指している²⁹⁸。

²⁹² <https://denvergov.org/Sunny>

²⁹³ <https://www.sf.gov/reports--december-2023--san-francisco-generative-ai-guidelines>

²⁹⁴ 同上

²⁹⁵ <https://consumer.ftc.gov/consumer-alerts/2023/11/announcing-ftcs-voice-cloning-challenge>

²⁹⁶ <https://www.ic3.gov/PSA/2024/PSA241203>

²⁹⁷ 同上

²⁹⁸ <https://about.fb.com/news/2024/04/metass-approach-to-labeling-ai-generated-content-and-manipulated-media/>