

AI生成画像の外部公開ガイドライン

AIを用いて生成した画像を社内外で活用・公開する際には、著作権や倫理面など様々な観点で注意が必要です。本ガイドラインでは、**著作権とライセンス、プラットフォーム別ポリシー、倫理的配慮、プライバシー・肖像権、商標・ブランド侵害のリスク、日本における法規や指針、そして主なリスク事例と対策**について、総合的に解説します。それぞれの項目で具体的なポイントや事例を挙げながら、AI生成画像を安全かつ適切に公開・利用するための指針を示します。

著作権とライセンス

- **AI生成画像の著作権:** 一般に、AIが自動生成した画像そのものには著作権が認められないと考えられています。日本の文化庁も「制作過程のすべてが生成AIによる場合、その生成物に著作権は認められない」としています。つまり**人間の創作的関与がない純粋なAI画像は、法律上“無権利”**となり、誰の著作物でもない状態です（事実上パブリックドメインに近い扱い）。そのため、生成者が独占的権利を主張するのは難しく、他者に無断使用されても著作権侵害で訴えることは困難です。ただし**人間が詳細なプロンプト指示や加工で創作性を付与した場合**には、著作物とみなされる余地があり得ます（創作意図と創作的寄与が認められるケース）¹。
- **他人の著作物を学習・模倣していないか:** AIは膨大な既存データを学習しているため、**生成結果が知らず知らず既存の著作物と酷似してしまうリスク**があります。とくに「○○の作風で」など特定アーティストのスタイルを真似る指示をすると、その作品の特徴を写し取った画像が出力される可能性があります。その場合、単なる「画風」の類似であれば著作権侵害には該当しません。文化庁の見解でも**画風・作風はアイデアに過ぎず保護対象ではない**ため、画風が似ているだけなら直ちに侵害にはならないとされています。しかし、**構図やキャラクターなど具体的表現まで既存作品と共通する場合は侵害となり得ます**。実際に中国では、生成AIが有名キャラクター「ウルトラマン」に酷似した画像を出力した件で著作権侵害が認定され、サービス提供者に賠償命令が出た事例があります。このように**AI生成物であっても既存キャラクターや作品の独創的表現を複製していれば違法となる可能性**があります。
- **商用利用とサービス規約:** 生成した画像を商用目的で使う際は、**利用した生成AIツールのライセンス条件を必ず確認**しましょう。サービスによっては無料プランでは商用利用不可だったり、有料契約が必要な場合があります。例えば**OpenAI (ChatGPT / DALL-E)**は有料契約者には生成物の商用利用を許諾していますが、**出力内容の正確性や権利侵害がないことは保証されない**ため、最終的な責任は利用者にあります。**Midjourney**などのサービスも有料プランでは商用利用可・著作権譲渡あり、無料プランでは生成物が公開ドメイン扱いになるなど条件が異なります。こうした規約を守らないと、後日ライセンス違反を問われるリスクがあります。また、**安全策として権利処理済みのデータで学習したAIツールを使うのも有効**です。例えば**Adobe Firefly**は学習データに著作権フリー素材のみを用いており、商用利用可能な**生成物に第三者の権利侵害があった場合には1素材あたり最大1万ドルの補償まで提供**しています。このような**IP保証付きサービス**を選ぶことで、万一のトラブルに備えることができます。
- **利用者側でのチェック義務:** AIサービス側は「学習データ由来の出力に関する保証はしない」ことがほとんどであり、**生成物に他人の権利侵害がないか確認する責任は利用者にあります**。商用利用前には、**出力画像に既存のキャラクターや企業ロゴが紛れ込んでいないか、著名な写真やイラストと酷似していないか**を入念にチェックしましょう。必要に応じて逆画像検索を使ったり、専門家に相談することも検討すべきです。また**可能であれば生成画像をそのまま使わず多少なりとも編集・加工して独**

自性を高めることで、他人の著作物と実質的に異なる表現にする配慮も有効です。企業での利用なら、社内ガイドラインでチェック体制を整備し、クリエイター任せにしない仕組みを作ることが望ましいとされています。

プラットフォーム別の規定

公開先のプラットフォームごとに、AI生成コンテンツに関するポリシーやルールにも違いがあります。主要プラットフォームの動向を押さえておきましょう。

YouTube

YouTubeではAI生成コンテンツ、とくに実写と見分けがつかない合成メディア（いわゆるディープフェイク）への対応を強化しています。2023年末にYouTubeは「現実的なAI合成コンテンツには開示が必要」との方針を打ち出し、2024年には具体策を実施しました。現在YouTubeでは:

- ・リアルに見えるAI生成・改変コンテンツはアップロード時に“Altered content（改変コンテンツ）”であることを自己申告する義務があります。アップロードフローで該当オプションを選択すると、動画の説明欄に「この動画にはAIによる合成コンテンツが含まれます」といったラベルが表示されます。例として、**実在の人物が発言していない言葉を言ったように見せる映像、実在の場所で実際には起きていない出来事のリアルなCG映像**などは必ず開示が必要です。
- ・明らかにフィクションと分かる合成（例: おとぎ話の世界や明確にCGと分かる加工）は開示不要ですが、**視聴者が現実と誤認しかねないものはすべて開示対象**です²。開示ルールに違反した場合、当初は注意喚起に留められますが、**悪質な場合は将来的に収益化停止などのペナルティも科す計画**が公表されています。
- ・自社のAIツール利用時の自動開示: YouTube Shorts向けの合成背景機能（Dream Screenなど）や公式AI機能を使った場合、プラットフォーム側で自動的に「AI使用」ラベルが付与されます。第三者のAIツールで生成した映像を投稿する場合は、クリエイター自身で上記の開示設定をする必要があります³。

このようにYouTubeは「**視聴者がそれを本物と思い込むリスク**」があるAIコンテンツには**透明性を確保**する方針です。動画説明欄の開示だけでなく、将来的には動画上に直接ラベル表示する取り組みも進めています。

X（旧Twitter）

X（旧称Twitter）でも、**合成・改変メディアの悪用に対するポリシー**が設けられています。Twitter社は2020年前後から「誤解を招く意図で改ざん・捏造されたメディア」への対策をルール化してきました。具体的には:

- ・他人を欺く目的で操作された画像・動画・音声の投稿は禁止されており、発見した場合は「**操作されたメディア**」という警告ラベルをツイートに表示したり、ユーザーがリツイート/いいねしようとする
と警告ポップアップを出す措置が取られます。例えば有名人の発言動画を勝手に合成した場合、そのツイートに「操作メディア」と明示される可能性があります。
- ・**深刻な被害のおそれがある場合は削除**: 合成メディアが原因で人々の生命や安全が脅かされたり、重大な被害に繋がりがねない場合、Twitterはその投稿自体を削除するとも明言しています。悪意あるディープフェイクやフェイクニュース拡散への強い牽制です。

- ・**パロディや明確なフィクションは対象外**: 政策上、政治的ミームやジョーク画像のような「明らかに風刺・パロディと分かるもの」は直ちに禁止とはしていません。しかしそれらも第三者を著しく誤解させる文脈では制限対象となり得ます。

要するにXでは、**合成画像による詐欺・デマの拡散防止**を目的にルールが敷かれていると言えます。「操作されたメディア」ラベルが付けば投稿の露出は抑えられ、悪質な場合はアカウント凍結等も起こり得ます。AI生成画像を投稿する際も、**他人を誤解させない表現かどうか**をよく考えることが求められます。

Instagram・Facebook (Meta)

Meta社（FacebookおよびInstagram）も2020年に**ディープフェイク動画の禁止方針**を打ち出しています。基本的な考え方は「**AI技術で人に虚偽の言動をさせるような合成動画は削除する**」というものです。主なルールは:

- ・**他者が実際に言っていないセリフを言ったように見せる合成動画、本物に見えるが完全に捏造されたリアル映像などは、発見次第プラットフォームから削除されます。**これはFacebook社が公式に発表したポリシーで、Instagramにも適用されています。例えば政治家が発言していない内容を喋っているフェイク動画などは即削除対象です。
- ・**パロディ・風刺の例外**: 政策では「パロディや風刺目的のコンテンツはこの限りでない」とされています。したがって、明らかにジョークと分かるディープフェイク（コメディ番組でのネタなど）は直ちに削除されません。ただしユーザーがそれを真実と誤認する恐れがある場合は、**第三者機関のファクトチェック**に回されます。ファクトチェッカーが「偽情報」と判定した場合、その投稿はフィードでの表示頻度が大幅に下げられ、閲覧前に警告が表示されます。広告で出そうとしても却下されます。
- ・**一般的なコミュニティ規範の適用**: そもそもInstagram/Facebookの利用規約上、**ポルノ・ヘイト・暴力扇動**などは禁止コンテンツです。これはAI画像にも当然適用されます。AI生成か実写かを問わず、**ヌードや過激な画像を投稿すれば従来通り削除**されます。要は**AIだから特別許されることはない**ということです。

以上から、Instagram等でAI画像を投稿する際は「**誤情報や有害情報の拡散につながらないか**」「**既存のコミュニティガイドラインに違反していないか**」をチェックする必要があります。最近ではMetaもAI生成コンテンツであることを明示する実験（例：画像に透かし埋め込み）を進めており、**ユーザー自身も投稿時に#AIartなどハッシュタグで自主的に明示する動き**が広がっています。

企業ウェブサイト・広告媒体

自社のウェブサイトや広告（パンフレット、SNS公式アカウント等）にAI生成画像を掲載する場合、明確な「プラットフォーム規約」はありません。しかし**社外向けに発信する以上、社会的な信頼や法的リスクに留意**する必要があります。以下のポイントに注意してください。

- ・**権利関係の明確化**: 企業が制作物にAI画像を使う場合、その**素材に第三者の権利侵害がないと説明できるか**が問われます。2023年には**海上保安庁が広報パンフレットにAI生成イラストを使用したところ、「著作権処理が不明」と批判を浴び、結局パンフレットを回収・廃棄する事態**になりました。このように公的機関・企業での利用では透明性が特に重視されます。利用する際は「どのツールで生成し、著作権はどうか」を関係者に説明できる状態が望ましいでしょう。また、素材提供サイトから入手したAI画像の場合は、そのサイトの利用規約（商用利用可か、AI生成物の投稿ポリシー等）も確認が必要です。

- ・**顧客やファンの感情に配慮:** 企業のマーケティングでAI画像を使うと、「手抜きでは?」「本物を使わないのは不誠実」といった反発が起きる場合があります。実際、大手企業でもAI画像利用が発覚して炎上した例があります。例えば花王「バブ」のパッケージイラストがAI生成ではないかと疑われ、SNSで議論が巻き起こったケースでは、企業は「人間が描いた」と説明しましたがユーザーの不信は完全には晴れませんでした。また回転寿司チェーンのスシローがSNS広告にAI生成の寿司画像を使用したところ、「本物の写真を使わないのはなぜ?」と批判が集まった例もあります。特に食べ物やファッションなど「本物らしさ」が大事な商品ではAI画像は逆効果になる可能性があります。顧客層の期待する価値とAI活用のバランスを考慮しましょう。
- ・**広告表示法との関係:** 日本では景品表示法という法律で、商品の実態以上に良く見せる誇大表示が禁じられています。AI画像で実物より美しいビジュアルを作り出し、それをあたかも実物写真かのように使うと、景表法違反（優良誤認表示）に問われるリスクがあります。例えば料理のボリュームや色味をAIで盛った画像、モデルの肌をAIで極端に美しく加工した画像を商品宣伝に使えば、「実際と異なる表示」と判断されかねません。特に食品・化粧品分野では取り締まりが厳しいので注意が必要です。商品写真や効果イメージにAI生成を用いる場合は、その旨を注記する（例：「イメージはAIによる合成です」）か、そもそも実物写真を使うなど、誤認させない工夫をしましょう。
- ・**自社ガイドラインの策定:** 複数の部門で生成AIを活用する場合、社内で統一の利用ガイドラインを作成することを強くお勧めします。ガイドラインには「利用できるAIサービスの範囲」「商用利用前のチェック項目」「公開時のクレジット表記ルール」「不適切な利用の禁止事項」などを盛り込みます。社内ルールを定め周知することで、現場ごとの判断ブレやリスク見落としを 방지、万一問題が起きた際の対処フローも明確になります。幸い日本ディープラーニング協会（JDLA）が企業向けに「生成AI利用ガイドライン」の雛形を公開しており、無料で参考にできます。自社の状況に合わせてテンプレートをカスタマイズし、社内規程化するとよいでしょう。

倫理的配慮（差別・偏見表現、AIである旨の明示 など）

AI生成画像の公開にあたっては、法律遵守だけでなく**社会倫理や価値観への配慮**も重要です。以下の点に留意してコンテンツの健全性を確保しましょう。

- ・**バイアスやステレオタイプの除去:** 生成AIは学習データに偏りがあると、人種・性別などに関する偏見を含む画像を生成する恐れがあります。例えば特定の職業イメージを描かせると男性ばかり生成される、肌の色で差別的な表現を含む、といった問題が起こり得ます。実際にAIチャットや画像生成で差別的・名誉毀損的アウトプットが出力された例も報告されています。**公開前に画像内容をよく確認し、特定の人種・民族・性別を不当に貶めたり固定観念を助長する表現がないか**チェックしましょう。必要に応じて人間が修正し、中立で多様性に配慮した表現に整えることが望ましいです。また、学習データ上避けられないバイアスがありそうなテーマ（歴史的事項や宗教的モチーフなど）については、AIの利用自体を慎重に判断することも必要です。
- ・**暴力的・性的表現の制御:** 当然ながら**公序良俗に反する画像**はAIであれ人間制作であれ公開すべきではありません。多くの生成AIサービス自体が利用規約でポルノや過激描写の生成を禁じています。万一そうした画像が生成されてしまった場合も、公開しない・即廃棄するのが賢明です。**特に性的なディープフェイク**（実在人物の顔をポルノ画像に合成したもの）は深刻な人権侵害です。日本でも未成年者の顔を合成した性的画像は児童ポルノと見なされ処罰対象になっています。**他者を性的に加工した画像の公開・提供は厳禁**であり、成人であっても人格権侵害や名誉毀損で民事・刑事上の責任を問われる可能性があります。
- ・**AI生成であることの明示:** コンテンツがAIによる合成である旨を積極的に開示することは、倫理的透明性の観点から推奨されます。法的義務がない場面でも、あえて明示することで誤解や混乱を 방지、受け手の信頼を損ねません。日本政府も新たなAI法制の方針として「AIの利用を透明化する（AI生成で

あることを適切に開示)」ことを掲げています。例えば企業サイトの素材写真にAI画像を使った場合、画像のキャプションや注釈として「※AIによるイメージ画像です」と記載する、といった工夫が考えられます。SNS投稿でもハッシュタグ「#AI生成」などを付与しておけば、ユーザーに变に勘繰られるリスクを減らせます。もっとも、**ネイティブ広告やエンタメ用途など敢えてAIと伏せたいケース**もあるかもしれません。しかしその場合でも、後から「実はAIでした」と判明すれば炎上につながりかねません。基本的には**見る側に誤認させないことを最優先に、開示すべきところではきちんと開示する姿勢が求められます**。なお、前述の通りYouTubeなど一部プラットフォームでは**現実的なAIコンテンツの開示が義務化**されていますので、それらに合わせる形で他媒体でも透明性を高めるとよいでしょう。

- ・**誤情報への加担防止**： AI画像はそのクオリティ次第で人々を欺き、社会に混乱を招く危険性があります。フェイクニュースサイトが偽写真を掲載したり、SNSでデマ画像が拡散された例も既に見られます。発信者として**自ら誤情報拡散の片棒を担がないよう十分注意**しましょう。例えば災害時にAI生成の偽画像をうっかり共有するとパニックを助長しかねません。特にジャーナリズムや公共性の高い文脈では、**AI画像は使わないか、使う場合も「これは想像図である」と明示する**など節度ある対応が必要です。政治家や有名人の写真をAI加工する場合も、その意図を明確にしておかないと**名誉毀損や選挙干渉と受け取られる**恐れがあります。要するに、**AI画像によって誰かを誤導・欺罔し得る状況は絶対に避けるべきです**。

プライバシーと肖像権

AIによって実在の人物に酷似した画像が生成できてしまうことから、**個人のプライバシー権や肖像権の侵害リスク**にも注意が必要です。

- ・**実在人物そっくりの画像**： たとえ有名人本人の写真を直接使っていないくても、**AIが作り出した顔が特定の有名人に瓜二つ**ということが起こり得ます。こうした画像を**本人非公認で公開・利用すると「肖像権」あるいは著名人の場合「パブリシティ権」の侵害**を問われる可能性があります。パブリシティ権とは、有名人の氏名・肖像を無断で商品宣伝等に使われない権利のことです。日本の判例でも有名人そっくり人形やロボットの展示がパブリシティ権侵害と認定された例があります。同様に、**AIで作った偽の有名人画像を広告やサムネイルに使う行為は極めてリスクです**。特定個人をモデルにしたような生成（〇〇に似た顔を作る等）は避けるのが賢明でしょう。どうしても著名人を題材にAI画像を作りたい場合は、私的利用に留め公開しないか、本人や権利管理事務所に許諾を取るべきです。
- ・**本人の同意なきAI合成**： 他者の写真を勝手にAIの入力に使って加工・生成する行為にも注意が必要です。それが公開されれば**プライバシー権や肖像権の侵害として民事で損害賠償請求**され得ます。特に性的な文脈での合成（リベンジポルノ的な使い方）は深刻な人格権侵害であり、先述のように**法規制の対象**です。また、AI音声合成で有名人の声を無断使用するのもトラブルになっています。実例として、歌手の高橋洋子さんは自身の声をAI模倣したコンテンツを使用するイベント出演を直前で辞退し、無断AI利用への問題提起を行いました。このケースは声ですが、画像でも同様に**他人の容姿を勝手に合成したコンテンツに本人が異議を唱えることが十分考えられます**。家族や友人であっても、了承なくAIで似顔絵を作りネット投稿するのは避けるのが無難です。
- ・**人物画像の商用利用の慎重さ**： 広告や販促物にAI人物画像を使う際は、その人物が架空でも**実在の誰かに似ていないか確認**しましょう。不特定多数が見る媒体で「あれは〇〇さんでは？」と誤認されるほど似ていると後で問題化する恐れがあります。特に有名人そっくりなら上記パブリシティ権の問題ですし、一般人でも本人が不快に思えばクレームになるかもしれません。どうしてもリスクを避けたいなら、**AIでなくモデル本人に登場してもらい適切に肖像権処理した写真を使うのが確実**です。AI人物を使う場合は**念のため生成時に特定人物の特徴を避けるプロンプト設定**をするとよいでしょう（例：「実在の人物には似せない」など）。また、一部サービスでは出力人物が実在しないことを証明する

「モデルリリース」的な証明書を発行する機能もあります。そうしたものを活用し、後から疑義を持たれた際に説明できるよう準備しておくことも考えられます。

- ・**顔認識や個人情報の流出**: AI画像に限りませんが、公開する画像内に写り込んだ人々のプライバシーにも配慮しましょう。街並み画像をAI加工した場合でも、そこに実在の通行人の顔が鮮明に残っていれば肖像権侵害の可能性があります。基本的に**本人が特定できる容姿や背景情報は、AI加工の有無にかかわらず許可なく公開すべきでない**という原則を忘れないでください。

商標・ブランド等の侵害リスク

AI生成画像は意図せず**既存のブランドロゴやキャラクターに酷似した要素**を含んでしまうことがあります。これは**商標権や著作権、意匠権などの知的財産権侵害**につながるため注意が必要です。

- ・**有名キャラクター・芸能人の画像**: 前述したとおり、著名なキャラクター（アニメ・ゲームの登場人物等）そっくりの画像は著作権侵害になり得ます。また、たとえ顔だけ別人でも**コスチュームやポーズが明らかに特定キャラを連想させる場合**も危険です。ファンアート程度なら黙認されることもありますが、公の場で使えば企業からクレームが来る可能性があります。「**〇〇風のヒーロー**」「**ジブリ風背景**」など有名IPを匂わせる生成は避けるのが無難でしょう。
- ・**企業ロゴ・商品デザイン**: AIが出力した街の風景に、偶然にも**実在企業のロゴに酷似したマーク**が写っている、といったケースも考えられます。公開してから指摘されると慌てることになります。商標権侵害の疑いを持たれぬよう、**画像内の看板やシンボルマークにも目を配り、有名ブランドのロゴやキャラクターが紛れ込んでいないか**チェックしましょう。例えば赤と白の〇〇社ロゴに似た図形、三日月マーク、有名スポーツチームのユニフォーム風デザインなど、権利物に似た要素は修正・削除すべきです。
- ・**商品パッケージの酷似**: デザイン分野では、AI生成物が既存商品に酷似し意匠権侵害を指摘されるリスクもゼロではありません。特に「有名な高級バッグのような画像を生成」など指示すると、本物と見間違ふような画像が出てくることがあります。このまま広告に使えばブランド側から法的措置を検討されるかもしれません。**他社の商品やマスコットを暗示する生成**は行わず、万一出力に含まれていた場合は使わないよう徹底しましょう。
- ・**訴訟例に学ぶ**: 商標・著作権侵害に発展したケースとしては、海外で**Getty Images社（ストックフォト大手）がStable Diffusion開発元を著作権侵害で提訴**した例が知られます。AIが出力した画像にGettyの透かしロゴ風の模様が残っていたことなどが発端でした（※このケースは学習段階の無断使用も問題視したもの）。また前述の**ウルトラマン画像訴訟（中国）**も、生成AI提供側が賠償責任を負うという厳しい判決でした。**ユーザー側も「知らなかった」では済まされない時代**になりつつあります。生成AIを使う企業も第三者の商標・著作物への配慮を怠らないようにしましょう。

日本における特有の法規やガイドライン

日本国内では、政府機関や業界団体が相次いで生成AI利用に関するガイドラインや見解を打ち出しています。法規制自体は2025年時点で欧州ほど厳格ではありませんが、「**責任あるAI利活用**」のための指針が多数公表されています。その主なものを紹介します。

- ・**総務省のガイドライン**: 総務省は通信・メディア政策の観点から「**生成AI利活用に関するガイドライン**」を公開しています。ここでは**生成AI利用時のトラブル対処、AI利用の基本概念、AI活用の10原則**などが解説されており、ビジネスで生成AIを導入する際に留意すべき事項が網羅されています。10原則には「人間中心の原則」「安全の確保」「透明性の確保」「知的財産尊重」などが含まれており、企業が社内ポリシーを作る際の参考になります。総務省はまた、生成AIを巡る課題整理やトラブル事

例の共有も進めており、利用者自身が自主的にチェックリストを作成するなど**リスク予防策を推奨**しています。

- ・**経済産業省のガイドライン**: 経産省は主に企業のAI導入を念頭に「**AI技術の社会実装指針**」を策定しています。ここでは**AIガバナンス（AIの適切な管理統治）**に向けた考え方が示されており、リスクベースでAIを評価・制御する方法が提案されています。生成AIに特化した部分では、**国際動向（EUのAI規則など）**や**ステークホルダーの懸念**を踏まえ、日本企業が取るべき対策が述べられています。例えば、著作権や個人情報の扱いに関する留意事項、AIアウトプットの品質管理などです。経産省指針は**企業が自主ルールを策定する際に非常に参照しやすい構成**になっています。
- ・**文化庁の見解（著作権）**: 文化庁は法律所管の立場から「**AIと著作権**」に関するQ&Aや資料を公表しています。これによれば、**AI生成物についても他人の著作物を無断で利用・複製すれば通常の著作権侵害になることが明確に**されています。一方で前述のように画風の模倣だけでは侵害ではないなど、グレーゾーンについても一定の考え方を示しています。さらに、**生成物それ自体の著作物性**についても触れられており、「完全自動の生成物は著作物でない」「人間が関与した場合、その関与の程度次第」という整理がなされています。文化庁のこれら見解は**現行法の範囲内でどうAIを扱うか**を示したもので、法律上の紛争が起きた際の判断材料になるため重要です。
- ・**文部科学省の指針（教育分野）**: 文科省は2023年、学校教育でのChatGPTなど生成AI活用に向けた暫定指針を発表しました。ここでは**子供たちがAIの回答を鵜呑みにしないよう指導することや、レポート課題等でのAI使用ルールを明確化すること**などが提言されています。学校現場向けですが、「AIが常に正しいとは限らない」「出力された情報の真偽をチェックする重要性」といったポイントは一般企業や個人にも通じる内容です。
- ・**民間団体のガイドライン**: 一般社団法人ディープラーニング協会（JDLA）や情報処理学会なども、**生成AI利用ガイドラインやガイドライン策定支援ツール**を公開しています。JDLAのテンプレートは企業が自社ルールを作る際のたたき台として有用で、**必要事項が網羅された雛形**になっています。他にも業界別では映像製作団体がAIのクレジット表記指針を出したり、クリエイター支援団体が著作権に配慮したAI活用の提言をまとめたりしています。さらに、**2025年には日本版の包括的なAI法（生成AIを含む）も施行予定**であり、その基本理念として「AIの発展を促しつつ危険な使い方は抑制する」という方針が示されています。同法では**コンテンツのAI利用の透明性確保や基本的人権の尊重**などが謳われる見込みで、企業・クリエイターもその精神に沿った対応が求められます。

総じて日本では、**強制力のある罰則よりもガイドラインによる自主的な安全確保**が重視される傾向です。その分、利用者一人ひとり・各企業のリテラシーが試されると言えます。「最新の指針に常に目を配り、自社のルールをアップデートし続けること」が重要です。

よくあるリスク事例と対策の教訓

最後に、実際に起きたAI生成画像絡みのトラブル事例をいくつか紹介し、それぞれから学べる教訓をまとめます。

1. **官公庁パンフレットへのAIイラスト使用事件** - （前述）海上保安庁が2023年、広報資料にAI生成のイラストを採用したところ、「著作権処理は大丈夫か？」との批判が集まり、急遽パンフレットを回収する事態になりました。 **教訓**: 公的な場では特に、AI画像の利用について事前に十分な権利確認と説明が必要。不安を与えれば信頼失墜につながるため、場合によってはAI素材の使用自体を控えることも検討すべき。
2. **著名アーティストの無断AI利用問題** - 人気歌手の高橋洋子さんは、自身の声をAIで合成したコンテンツを使うイベントへの出演を辞退しました。この出来事は、許諾なきAI利用への異議表明として話題

に。 **教訓：** 人の声や姿形といった人格的な要素を勝手にAIで再現して利用することは、当人の人格権・肖像権を侵害し得る重大な問題。芸能人などの素材を扱う際は必ず許諾を取り、少しでも不安があれば使用を控える。

3. **写真コンテストでAI作品が入賞** – ある国際写真コンテストで、応募者がAI生成画像を人間の撮影作品として紛れ込ませ、受賞してしまった例があります。後日それがAI作品と判明し、賞は剥奪、主催者・審査員も大きな批判にさらされました。 **教訓：** コンテストや展示会など創作性を競う場では**AI産物は原則参加不可**か、参加する場合も**事前申告が必須**であるべきです。透明性なくAI作品を混入させる行為は、クリエイターコミュニティの信頼を損ね、自身も信用を失います。
4. **大手企業の商品デザイン炎上** – 入浴剤「バブ」のパッケージイラストがAI生成ではとSNS上で疑惑が広がりました。企業は人手によるイラストと説明しましたが、「AI風だ」「説明が本当かわからない」と議論は沈静化せず、企業イメージに影を落としました。 **教訓：** **大企業ほどAI利用の透明性が求められる**という指摘がされています。仮に人手であっても「AIでは？」と疑われれば炎上し得る時代です。デザインの過程も含めた情報開示やユーザーとの対話を厭わず、真摯に対応することがブランド保護につながります。
5. **飲食チェーンのAI料理画像トラブル** – 寿司チェーン「スシロー」がTwitterに投稿した握り寿司の画像が実はAI生成で、「本物を使えばいいのに」「フェイク写真で宣伝するなんて」と批判が巻き起こりました。寿司の品質そのものへの不信には繋がらなかったものの、「本物志向」の業界では不評という結果に。 **教訓：** **実在の商品や料理は、できる限り実物写真を使うのが鉄則**です。素材撮影コストを惜しんでAIに頼ると、かえって信用を損ないかねません。どうしてもAI画像を使う場合、「演出上のイメージです」と断るなどフォローが必要でしょう。
6. **文化イベントの偽写真拡散** – 銀座で開催された着物イベントの様子として、第三者がAI生成した画像をSNSに投稿し、人々がそれを実際の写真と信じて拡散してしまう出来事がありました。後に偽物と判明し、「日本文化の表現を勝手に歪めた」として議論に発展。 **教訓：** **SNS上ではAI偽画像が本物同然に流通する危険**があります。イベント主催者側も公式アカウントで誤情報を訂正するなどの対策が必要ですし、ユーザー側も容易に信じて共有しないリテラシーが求められます。**発信者としても、AI合成と実写が混同されないよう細心の注意を払うべき**です。
7. **元米大統領の逮捕フェイク画像** – 2023年、トランプ元大統領が逮捕されるシーンを描いたAI画像が大量に生成・拡散され、あたかも現実の報道写真のようにSNS上を駆け巡りました。一部では名誉毀損の声も上がり、フェイク拡散の深刻さを示す例となりました。 **教訓：** **政治家や公人に関するAIフェイクは社会的影響が極めて大きく、訴訟リスクも高い**ことを肝に銘じましょう。合成画像が生む誤解は社会混乱や人格権侵害につながります。公的な話題でAI画像を扱う際は特に厳重な配慮が必要です（基本的には使わない方が無難です）。

以上の事例から総合的に言える対策は、「**事前のチェック**」「**透明な開示**」「**無用な使用を避ける慎重さ**」に尽きます。思わぬところに法的・社会的リスクが潜んでいることを常に意識し、**公開前にリスクを洗い出してつぶすことが安全なAI活用への近道**です。社内での教育や専門家との連携も取り入れながら、生成AI画像を賢く使いこなしましょう。

参考文献・情報源： 本ガイドラインの内容は、総務省・文化庁など公的機関による資料や、大手プラットフォームの公式ポリシー発表、そしてAI利活用に関する専門家の解説記事等をもとにまとめています。各節末に示したとおり、具体的なケースや指針の引用元を記載しておりますので、詳細はそちら（各【+】内の出典）も参照してください。

① 生成AIの著作権侵害事例6選！著作権侵害を回避する方法も解説 | SHIFT AI TIMES

<https://shift-ai.co.jp/blog/5514/>

② ③ Disclosing use of altered or synthetic content - Android - YouTube Help

<https://support.google.com/youtube/answer/14328491?hl=en&co=GENIE.Platform%3DAndroid>