

Time for a check-up:
最新調査レポート

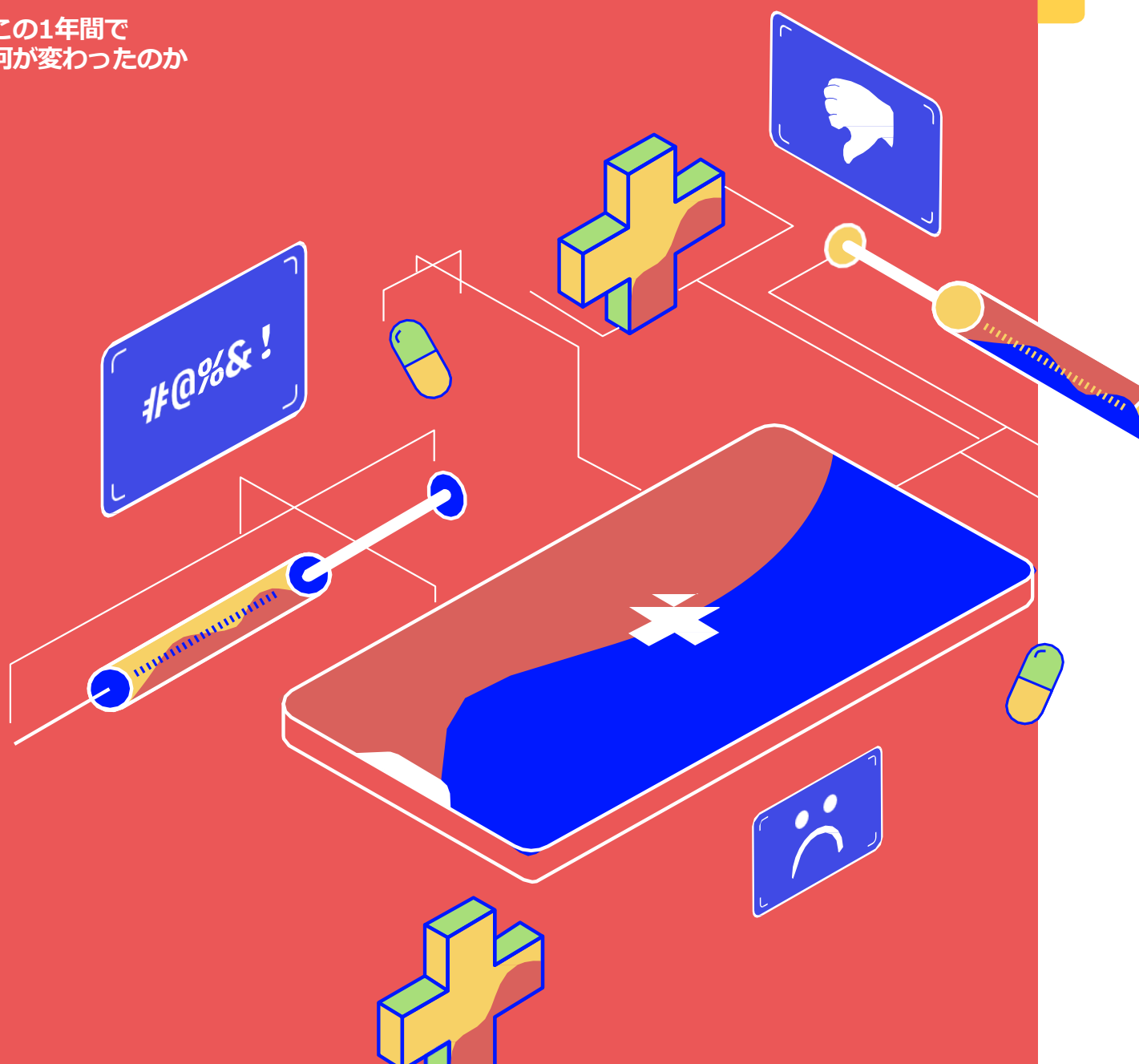
The brand safety

ブランドセーフティの現状

CRISIS

One year later

この1年間で
何が変わったのか



目次

ブランドセーフティは、 企業自身が直接取り組むべき 課題である。

p. 03

“企業のメディア担当者の中には、事実上「ブランドセーフティオフィサー（ブランドセーフティの専任者）」ともいえる業務を担う方が多く存在する。”

メディアプラットフォーム側の 努力が成果を上げ始めている。

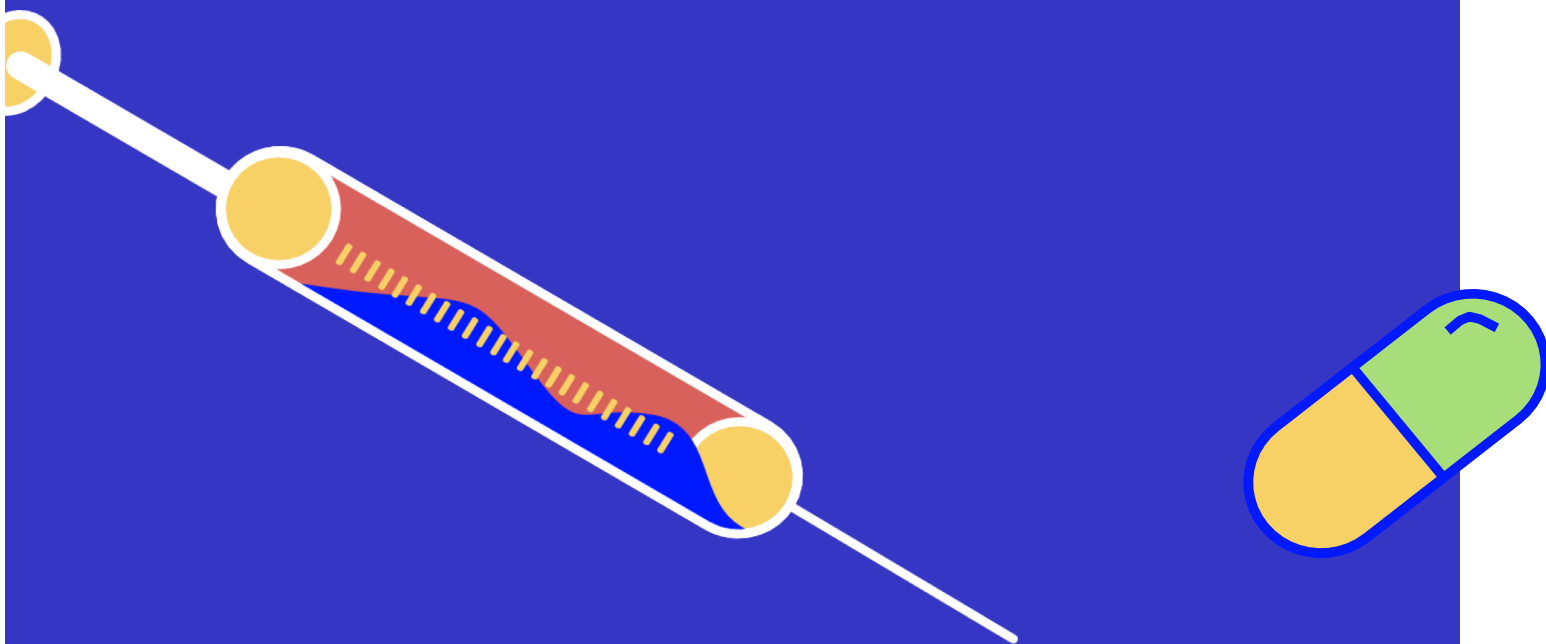
p. 04

“ソーシャルプラットフォームは、ブランドセーフティ問題を真摯に捉え、問題があることを認め、解決に努めている。”

マーケターは、優良なオー ディエンスを自ら遠ざけて しまっている。

p. 09

“現在、市場では過剰な対応が行われている。しかしそれはマーケターが悪いのではなく、ツールやシステムやエコシステムの問題で、ブランドが過剰な対応を行わざるを得ないほどのリスクが生み出されてしまっているからである。”



ブランド セーフティ の現状は、 —

前年と比べてどう変化したのか。

2017年後半、ブランドセーフティは危機的状況にありました。ブランドはパニックを起こし、主要なプラットフォームへの広告を取りやめ、ソーシャルメディアの幹部に怒りをぶちまけ、エージェンシーにはとにかく解決策を見出して欲しいと泣きつくような状況でした。日々の報道がこれを助長し、物議を醸すインフルエンサーは誠実なフォロワーを攻撃し、個人情報晒され、暴力・アダルト・フェイクニュース等に隣接した広告掲載にブランドは翻弄されていました。

2017年に実施した調査では、90%ものデジタルメディア専門家がブランドセーフティは深刻な問題であると回答しました。実際、依然としてエコシステムは危険をはらんでいます。

しかし、ブランドセーフティにかかる報道のあまりにも多くが、重要な視点を欠いていました。この1年間の間に多くの有効な解決策が開発され、実際に活用されてきました。画像解析や自然言語解析による文脈解析が、有効なテクノロジーの筆頭に挙げられます。

マーケターは、こうした解決策は全般として効果を上げていると捉えています。しかし同時に、ブラックリストやホワイトリストといった予防的な解決策の中には過度に制限的なものもあり、リスクから身を守ることではできても、ブランドにとって好ましい消費者さえも遠ざけてしまっていることをマーケターは懸念しています。



FOR Media Buyers

DUE TO Platform & partners risks

PRIMARY Brand RX

この1年間でブランド セーフティの現状は 変わったか。

マーケターやソーシャルプラットフォームは、暴力・アダルト・過激な思想などリスクの高い文脈でブランドが露出されることの危険性をいち早く認識しています。ブランド毀損リスクが伝染病のように広まったのはもう1年以上前のことですが、以降マーケターはさまざまな手段を講じてブランド毀損リスクの排除に努めてきました。ブラックリスト、ホワイトリスト、第三者による測定、ブランドセーフティの専門家やAIを活用したテクノロジーの導入などが代表例になります。

こうした「治療」は総じて効果を上げてきましたが、中には、オーディエンスプールの縮小や、ターゲティングの弱化など、新たなリスクをもたらしてしまう対策も見られました。

AIを活用したテクノロジー、雇用の変化、ソーシャルメディア環境の健全化などが功を奏し、「症状」は改善されブランドの安全性を損なう露出は減りました。しかし「症状」そのものは依然として残っており、オーディエンスプールの縮小やターゲティングの弱化といった「副作用」が残っているのもまた現状です。

SIGNED

GumGumはDigidayと共同で、2018年11月にブランド・広告代理店・パブリッシャーに在籍する274人の業界エキスパートに調査を行いました。その結果、殆どのマーケターが、主要ソーシャルプラットフォームが2017年以降行取り組んできたブランドセーフティ対策を評価していることがわかりました。多くの企業は自社内にブランドセーフティ対策を導入し、専門の担当者の採用も行っています。

しかし脆弱さは依然として残っており、多くの「治療法」が深刻な「副作用」を生んでいます。マーケターは予防策としてどのような対策を講じているのか、またそうした対策は効果を上げているのか、検証を行いました。

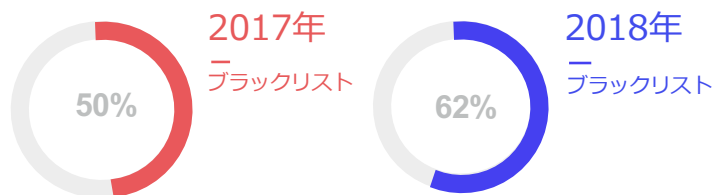
ブランドセーフティはマーケティングにおいて重大な問題か？



ブランド毀損リスクとして最も代表的なものは？



ブランドセーフティ対策に最も有効な手段は？



ホワイトリストやブラックリストによって意図する消費者へのリーチが難しくなったか？



ブランドセーフティ対策として画像解析を活用しているか？





ブランドセーフティは、 企業自身が直接取り組む べき課題である。

「伝染病」の広がり 企業の対応

ブランド毀損は、新しいリスク要因ではありません。2017年3月に、VerizonやL'Orealといった有名ブランドの広告が、ISISやネオナチ関連のYouTube動画周辺に掲載される問題が注目を浴びるまでは、あまり議論が行われなかっただけです。

多くのブランドがこうしたブランド毀損リスクを恐れ、広告出稿を実際に取りやめました。しかし一方で、多くのブランドが広告出稿を継続、または再開しています。

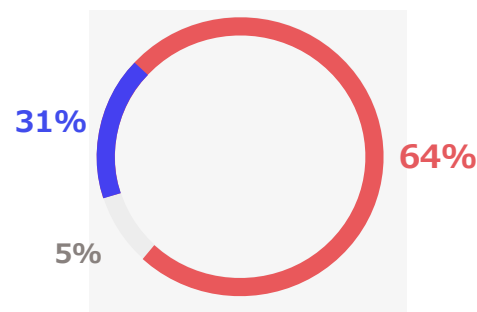
Facebook、Google Search、Twitterなど、その他プラットフォームについても同じ状況でした。「伝染病」があまりにも広がったため、リスクがあることを知りながらも、手を付けられずにいる傾向がありました。ここ1年間、マーケターの多くが自ら対処できる問題については然るべき対応を行ってきました。しかし、必ずしも新たなリソースを導入しているわけではなく、既存のリソースをこの問題の対処に充てているというのが実情です。

「当社クライアントのメディア担当者の多くは、事実上『ブランドセーフティオフィサー（ブランドセーフティの専任者）』としての業務を担っている方が多く存在します」とGroupMのJoe Barone氏（Managing Partner of Brand Safety）は語ります。Barone氏によると、こうした専門家は社内の既存マーケティングチームの中から抜擢されています。

一方で、Bank of America、IPG Mediabrands、GroupMといったブランドやエージェンシーは、専任の「ブランドセーフティオフィサー」を外部から新たに採用しており、各社様々なアプローチで課題への対応を進めています。いずれにせよ、専門の人材を配置するようになったということは、問題の解決に向けて企業が真剣に動いていることを示唆しています。

こうした専門の人材配置によって、どのくらい効果が生まれているのでしょうか。実際、マーケターの64%がこうした専門の人材は「多少は」役立っていると考え、31%が「それなりに」役立っていると感じています。「非常に」役立っていると感じているのは、全体のわずか5%にとどまりました。

- 多少は役に立っている
- それなりに役に立っている
- 非常に役に立っている



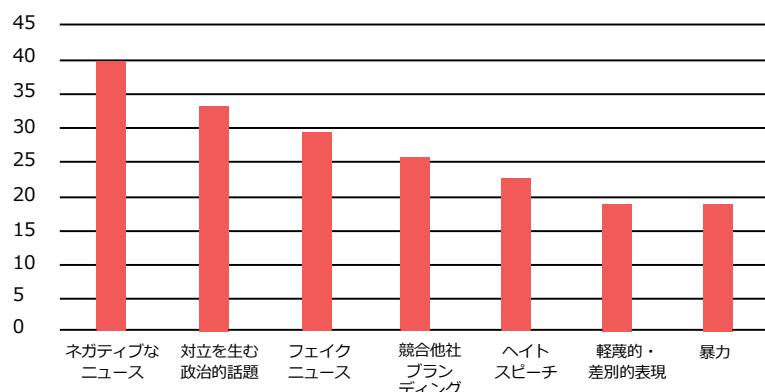
調査方法について

GumGumブランドセーフティ調査では、2018年10月27日～11月25日の期間中、米国内の広告主・広告代理店・媒体社・テクノロジー企業に勤務する274人の業界エキスパートに対し、インタビュー調査を実施。

メディアプラット フォーム側の努力が 成果を上げ始めている。



2017年

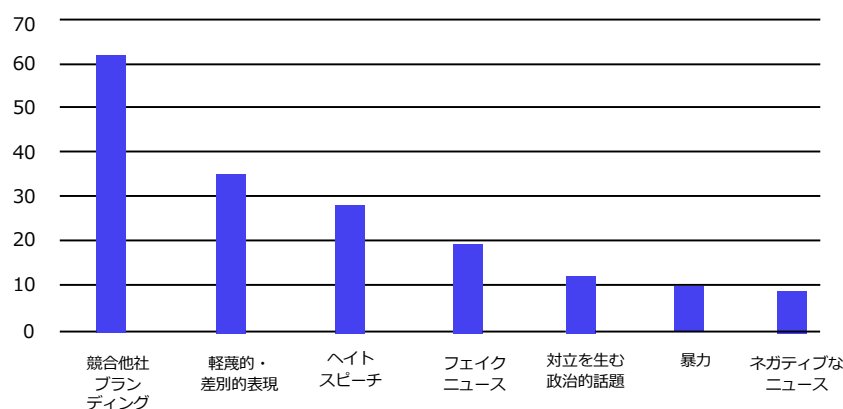


脅威のベクトル

マーケターが最も脅威だと感じるコンテンツのタイプは、この1年間で大きく変化しました。2016年の米大統領選挙から2年を経て、ニュース関連コンテンツ周辺でのブランド露出に対する懸念は弱まりました。

一方で、63%ものマーケターが、競合他社ブランド周辺での露出を恐れており、これが最もブランドの安全性を脅かすものであると考えています。前年度調査では28%だったのに対し、劇的な変化が見られました。

2018年



「広告取引の際、Fordの広告をGMの広告の隣に掲載したり、McDonald'sの広告をBurger Kingの広告の隣に掲載したりすることは無く、各ブランドがそれぞれ際立つように配慮を行っています」と動画ニュースサイトNewsyのCEO、Blake Sabatinelli氏は語ります。

「従来型のテレビメディアで行われていたのと同じことです」とSabatinelli氏は補足しています。競合他社とグルーピングされてしまうことに対する不安は今に始まったものではなく、古く以前から存在するものです。ただ、今ほどこの懸念がクローズアップされたことはない、ということが言えます。

デジタルマーケティング環境においては、ビジュアルコンテンツを取り扱うメディアが日増しに増えています。各ブランドは自社のロゴやその他のビジュアル資産が競合ブランドと同じスクリーンで表示されることで、自社ブランドの価値が中和されてしまったり、逆に競合ブランドにメリットを与えてしまうことが無いよう、ますます神経質になっています。広告主は膨大な資金を投入し、あらゆるデジタルメディアやソーシャルメディアで画像や動画を活用しています。1つの間違いが起こることで、何千回、何百万回もユーザーの目に触れてしまうリスクが生じるのです。

メディアで使用するビジュアルコンテンツにより多くの予算が使われるようになった現在、競合ブランド周辺での自社ブランド露出に過敏になるのは自然な流れであると言えます。

にもかかわらず、たった1年の間に競合他社ブランディングへの懸念が28%から63%に増加したことは、やはり驚きに値します。この理由の1つとして、他のコンテンツに関するブランドリスクの懸念が相対的に減少したことが考えられます。

こうした変化の背景には、ソーシャルメディアプラットフォーム側が課題解決のために行ってきた多大な努力があります。実際、マーケターの多くは、プラットフォーム側での自助努力の成果を認めています。

迅速な治療

FacebookやGoogleの過失が連日のように報道される中、こうしたプラットフォーム側での努力は見過ごされがちです。しかし、ブランドセーフティに関する業界の立ち位置は明確です。ブランドリスクという「伝染病」に対し業界として効果的な処置が行われ、FacebookやGoogleは其中で大きな役割を果たしています。

実際、調査対象者の多くが、主要ソーシャルメディアのブランドセーフティ対策への努力を高く評価しています。



「ソーシャルプラットフォームは間違いなくブランドセーフティ対策に真剣に取り組んでいます。課題があることを認め、処置を行っているのです」とBarone氏は語っています。

例えばFacebookでは、マーケターがインスタント記事やインストリーム動画のどこに自社広告が掲載されるかを把握することを可能にしており、特定のコンテンツカテゴリーやパブリッシャーへの掲載を防ぐことができるようにしています。またYouTubeは、マーケターが簡単にホワイトリストを設定し、事前に選択された信頼できるコンテンツグループ内だけにのみ広告を掲載できるようにしています。

更に重要な動きとして、Facebook・YouTube・Twitterといったプラットフォームが、第三者による測定評価のために積極的に情報開示を行うようになったことが挙げられます。これにより、エージェンシーやベンダーが、プラットフォームのどこに広告が掲載されるかを簡単に監視できるようになりました。

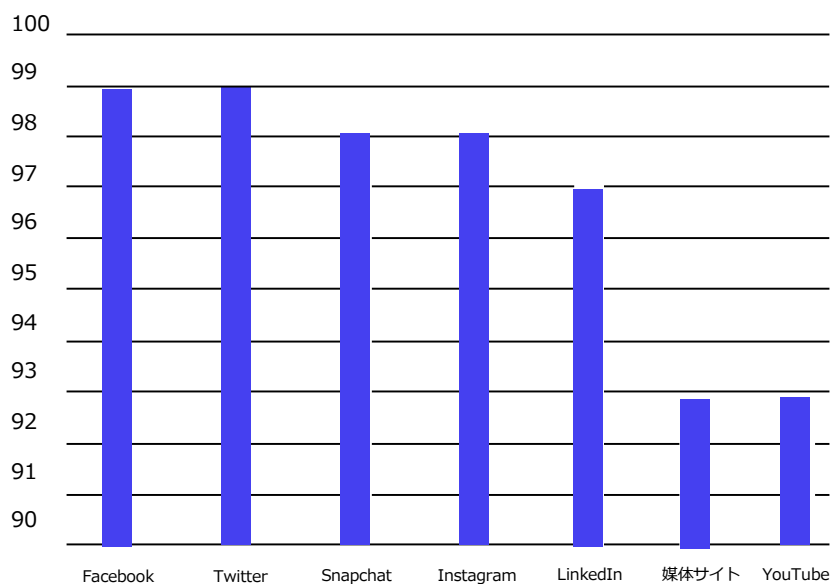
「やるべきことはまだまだあります。ただ、クライアントの大半はこの動きを見ており、デジタルメディアの信頼回復に対する強いコミットメントを目の当たりにしています」とBarone氏は語っています。

パブリッシャーも同様の見解を示しており、Sabatinelli氏は次のように述べています。「ソーシャルメディアで今起きていること、今後18～24カ月に起こるであろうことは、プラットフォーム側がより大きな説明責任や報道の責任を果たそうとしていることを示しています。」



過去12カ月にわたる
ブランドセーフティ状況の
改善について、どのプラット
フォームが適切な努力を
行ってきたと考えますか？

2018年



高い健全性が 認められたTwitter

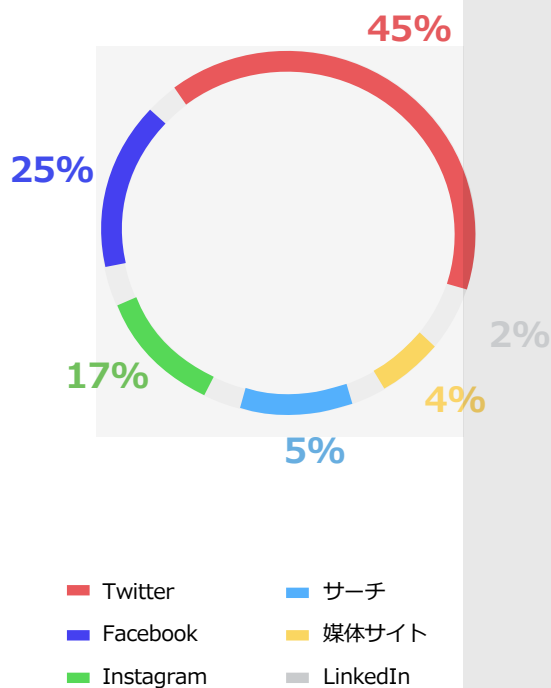
特にTwitterは、ブランドセーフティに対するアプローチを転換し、そのプロセスの中で信頼性を高めることに成功しました。

2017年の調査では、Twitterを「最もブランドセーフティに配慮したプラットフォーム」と答えた専門家はわずか1%でした。2018年にはその割合が45%となり、非常に多くの専門家の支持を得ました。

昨年の調査では、LinkedInがトップでした。今回の結果は、LinkedInというキャリアプラットフォームの安全性が下がったわけではなく、Twitterの安全性が非常に高く評価されたことを示しています。実際、LinkedInを「ブランドセーフティに問題があるプラットフォーム」と挙げた回答者はわずか4%でした。

今回のTwitterの結果は大きなトレンドを象徴していると言えます。この1年間、ソーシャルプラットフォームは非常に多くの時間をかけて「InfoWars」などの物議を醸したアカウントを閉鎖したり、暴力やアダルトといったブランドセーフティを脅かすコンテンツをプラットフォームから排除してきました。例えばTumblrが重い腰を挙げてようやくポルノを禁止したことは大きな話題となりました。面白くないと感じるユーザーもいるでしょうが、マーケターはこの動きを歓迎しています。

次の中で最も
ブランドセーフティに
優れたプラットフォーム
を挙げてください。



メディア擁護団体Sleeping Giantsの創始者Matt Rivitz氏は次のように指摘しています。

「ソーシャルプラットフォームは今まですべてを容認していました。例えば（白人至上主義者として知られる）Richard SpencerがTwitterのユーザーとして何の制限も受けずに発言することができていました。人種差別主義者であるにも関わらず、彼はTwitter・Facebook・YouTubeを自由に使い、プラットフォーム側は何も対策を行っていませんでした。この状況を受けて広告主側から『何か対策が講じられるまでは広告を取りやめる』という通達があったわけです。」

Twitterは第三者による測定評価を追加導入したほか、自動ツイートの「ボットアカウント」や不正なアカウントを排除するよう努力を重ねてきました。「不適切なコンテンツに関連付けられた広告露出問題を排除するよう、Twitterは非常に大きな努力を行ってきました。」とBarone氏は話します。

「TwitterはGoogleやFacebookと比べると規模の小さいプラットフォームです。巨大プラットフォームが批判の矢面に立たされているため、比較的規模の小さなプラットフォームがその陰で一息つくようなこともあります。」

一連の対策は自然発生的に起きたものではなく、2017年に「ブランドセーフティに問題のあるプラットフォーム」としてTwitterを挙げたマーケターからのプレッシャーによって対策が講じられるようになったのです。

「優良なプラットフォームは広告によって運営されています」とRivitz氏は話しています。「したがって収益を上げることは彼らにとって非常に重要です。外部からの世論の高まりと、何が自社のビジネスにとって正しいかという内部での議論の両方があるのはじめて、変わることが可能なのです。」

Twitterは市場の声に耳を傾けました。2018年には好調な業績を受けて株価が急上昇しましたが、これは全くの偶然ではないはずです。



マーケターは、 優良なオーディエンスを 自ら遠ざけてしまっている。

パブリッシャーへの 猜疑心

誰もがソーシャルメディアの動向に注目していますが、マーケターは他へも目を向けています。今回の調査で、全プラットフォームの中で最もブランドリスクが高いのは、独立系のウェブサイトであると考えられるマーケターが多いことがわかりました。45%ものマーケターが、パブリッシャーサイトにおけるブランドリスクが最も高いと考えており、更に65%がパブリッシャーサイト上でのディスプレイ広告が最もリスクの高い広告出稿であると捉えています。

この割合は、ソーシャルメディアの中で最も安全性が低いと考えられているFacebookでも26%で、大差をつけて2番目となっています。ブランドセーフティが問題視されるきっかけとなったYouTubeも、11%で3番目となっています。

マーケターにとって、脅威のベクトルが大きく変化しているのは明らかです。ほんの1年前には、34%ものマーケターがFacebookを最もブランドリスクの高いプラットフォームだと回答していました。当時、パブリッシャーサイトのリスクが最も高いとしたマーケターは27%で2番目、YouTubeが15%で3番目でした。



「過剰治療」というリスク

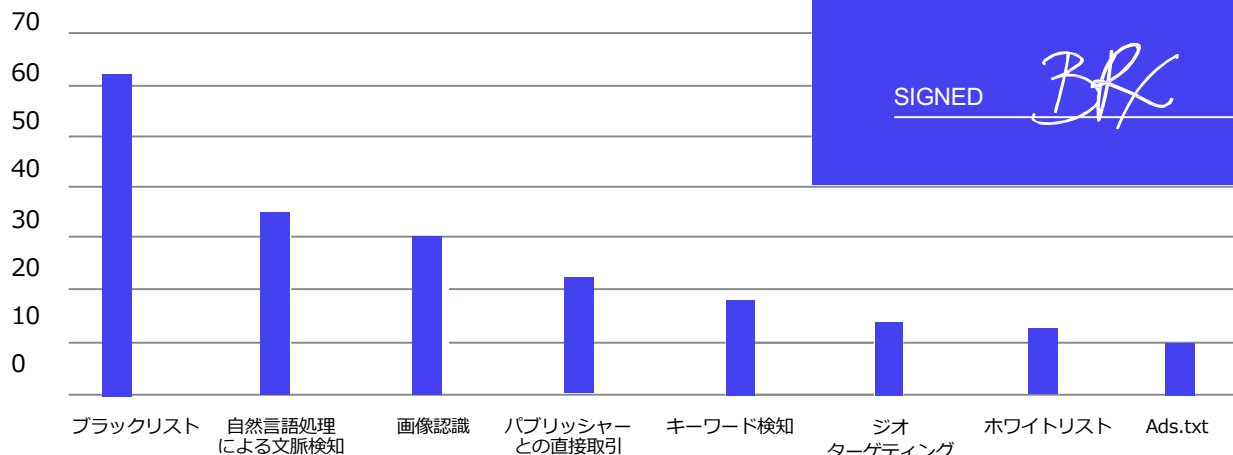
マーケターはブランドリスクという「病氣」を治そうと、さまざまな方法に救いを求めてきました。現在行われている手法を見ると、特定のパブリッシャーへの懸念が増加していることがわかります。この1年で、信頼できるパブリッシャーとの直接取引によってブランドリスクを未然に防ごうとしているマーケターの割合は55%へと増えました。これは他の手法よりもはるかに高い割合で、前年度の39%と比べても大きく増加しています。

また、効果面においては他の手法を高く評価するマーケターが多く見られました。広告掲載を避けたいサイトをリストアップして配信を制御する「ブラックリスト」の効果を認めるマーケターは実に62%にも上り、前年度の50%から大きく増加しました。

しかし一方で、大きな短所があるのもまた事実です。ブラックリストは事後対応として作成されることが多く、実際にリスクの高いブランド露出があって初めて、そのサイトをブラックリストに加えることもあります。この場合、既にブランドが多くのダメージを受けてしまっていることになるため、やはり予防的措置が最善の策であることは言うまでもありません。

ブラックリストが予防的に使用されたとしても、優良なオーディエンスを誤って遠ざけてしまう可能性があることは否めません。サイト内にブランドにとって理想的な内容が含まれている場合、本来はそれ以外の部分を単純に回避すれば良いだけのはずです。「ブラックリストは非常に有効な手法ですが、知らないことは知らないままであるという問題も抱えています」とBarone氏は指摘しています。

過去1年間で、どのブランドセーフティ対策が最も効果的でしたか（2つ選択）？



FOR Media Buyers

DUE TO Platform & partners

PRIMAR Brand RX

Y

過剰投薬に なっていませんか？

ブラックリストやホワイトリストが有効な措置であることは確かです。しかし使いすぎると、ブランドリスクだけでなく、優良なオーディエンスもブロックしてしまうことになりかねません。

**治療措置が、問題よりも悪い結果をもたらす
べきではありません。**

大多数のブランドが、ホワイトリストやブラックリストが初期症状には効果的であると考えていますが、一方で深刻な副作用も引き起こされてきました。「ブラックリストの問題は、知らないことは知らないままであるということです」とGroupMのJoe Barone氏 (Managing Partner of Brand Safety) は指摘しています。「ホワイトリストは広く使われていますが、脂肪を取り除こうとして肉の部分まで切り取ってしまうこともあるのです。」

ブランドの安全を脅かす露出は減りましたが、同時に広告出稿による効果も薄れてしまっている可能性があります。ブランドリスクから身を守り、制約の多いホワイトリストやブラックリストへの依存を緩和するためにも、画像スクリーニングや自然言語処理による文脈検知などでターゲティング機能を補うことを推奨します。

SIGNED

ホワイトリストは、ブランドが事前承認した安全なサイトのみに広告を掲載する手法で、マーケターの11%がこの手法を採用しています。しかし、広告が掲載されるのが少数のサイトのみに限定され、決まったサイト以外の優良な環境から自らを締め出してしまうというリスクがあります。「脂肪を取り除こうとして肉の部分まで切り取ってしまうこともあるのです」とBarone氏は指摘しています。

マーケターは、ブランドが守られているという点については自信を深めましたが、望ましいオーディエンスにリーチできているかについては確信を持てなくなってきています。実際、ブラックリストやホワイトリストなどのブランドセーフティ手法を使うことで、特定のユーザー層へのリーチが難しくなったと回答したマーケターは69%にも上ります。前年度の回答が30%だったことを考えると、これは劇的な変化であると言えます。

「市場では過剰な対応が行われてしまっているようです」とVox MediaのChief Revenue OfficerであるRyan Pauley氏は述べています。「しかしそれはマーケターが悪いのではなく、ツールやシステムやエコシステムの問題で、ブランドが過剰な対応を行わざるを得ないほどのリスクが生み出されてしまっているからなのでしょう。」

Pauley氏は、新たなテクノロジーがこの問題を解決できると考えています。例えばVox Mediaでは「カンバセーショナル・インテリジェンス (Conversational Intelligence)」というツールを活用しています。これは、自然言語処理と機械学習により、トーンや文脈の評価を行うものになります。

Pauley氏は次のように説明しています。「従来は下位5%のパブリッシャーから手を引くことに重点が置かれており、ブランドセーフティを取り入れることの意義については十分に検討されていませんでした。2019年には、『どのような文脈か』を問う必要が出てくるでしょう。ページ上にどのようなコンテンツがあり、誰がそれを閲覧しており、その読み手は何を考えているのか。さらに、文脈を捉えてブランドセーフティを実現するためにどんな新しいツールが使えるのか、といったことです。」



FOR Media Buyers

DUE TO Platform & partners

PRIMAR Brand RX

¥

ニュース関連の脅威のベクトル

「フェイクニュース」は依然として極めて高いブランドリスクですが、ブランドの心配は以前よりも減りました。ソーシャルプラットフォームやパブリッシャーは、この1年間にフェイクニュース問題に対処するために様々な対策を講じてきました。InfoWarsといった有害なアカウントは閉鎖されました。エージェンシーや第三者計測機関のエキスパートやブランドは、アカウントを検証し広告がどこに出稿されているかを見ることができるようになりました。ブランドや広告主はBreitbartといった物議を醸すアカウントをブラックリストに加えました。

パートナーを知る

広告主の多くは、特に正確で信頼できると考えるニュースパブリッシャーと緊密な関係を構築してきました。「二度目のチャンスはもはやありません」とNewsyのCEO、Blake Sabatinelli氏は指摘しています。「事実について誤った報道を行えば、それまでです。あらゆるものを、何度も何度も吟味することが当社の仕事なのです。」

回避が全てではない

世の中にはニュースが溢れており、すべてが心温まるものばかりではありません。対立を引き起こす政治的課題・暴力・企業の不正行為といったニュースは、ブランドにとって望ましいものではありません。しかし、こうしたニュースはユーザーの関心を集めてしまいます。「ブランドはどのような話題と自社を関連付けたいのか、自らが決断する必要があります」とVox MediaのChief Revenue OfficerであるRyan Pauley氏は説明します。「しかし、こうした話題を回避するのはますます難しくなっています。避けるべき話題を増やしてしまうと、認知の幅が極端に狭まることになります。」

SIGNED

ブランドセーフティの処方箋

パブリッシャーとの直接取引

信用のあるパブリッシャーから直接広告枠のバイイングを行うことで、ブランドセーフティへの懸念を早い段階で払拭することができます。さらに、直接取引を行うことで、ブランドセーフティへの懸念が生じた際に容易に対応依頼を行うことができるのも、ブランドにとってのメリットです。今回の調査では、マーケターの55%が予防策としてパブリッシャーとの直接取引に頼っており、54%が実際にブランドセーフティ問題に直面した後でこうした直接取引に頼るようになったと回答しています。前年度の調査では、同質問に対する回答はそれぞれ39%と34%でした。

処方箋

直接取引は、メディアバイイングの健全性と品質の担保に効果があり、即効性のある処置。

考えられる副作用

ハードニュース周辺でのブランド露出の可能性、割増価格設定

ブラックリスト

ブラックリストは、一般的にエージェンシーが作成しDSPが実装するもので、広告配信先として相応しくないと判断される特定のパブリッシャーサイトを除外するものです。マーケターの29%が予防策としてブラックリストを導入しており、28%が是正措置として導入したと回答しています。62%ものマーケターが、ブランド毀損リスクを最小限に抑えるためにブラックリストは最も有効であると答えています。

処方箋

ブラックリストは、ブランドにとって不利益な環境での露出を回避するために有効な手段。

考えられる副作用

オーディエンスプールの縮小、ターゲティング精度の低下

画像認識（コンピュータビジョン）

画像認識（コンピュータビジョン）はAIを駆使したテクノロジーで、ニューラルネットワークを使ってビジュアルコンテンツを識別し分類することが可能です。ブランド毀損リスクの高いコンテンツを検出し、リスクの高い環境を自動的に避けることができるのが特徴です。パブリッシャーやソーシャルメディアエコシステムには無数の画像や動画コンテンツが偏在しているにもかかわらず、画像認識を予防策として活用しているマーケターはわずか21%で、実際にブランドセーフティ問題に直面した後で導入したマーケターも16%に留まっています。前年度の同質問に対する回答はそれぞれ12%と10%で、ここ1年間で増加傾向にあります。

処方箋

ブランド毀損リスクのある画像や動画のモニタリングを行うことで、安全な環境下で見込み顧客へリーチすることが可能。

考えられる副作用

特になし

ブランドセーフティの処方箋

自然言語処理による文脈検知

自然言語処理による文脈検知は、画像識別と併用されることが多く、ブランド毀損リスクの高い文字情報を検知して回避するものです。21%のマーケターが予防策として導入しており、17%がブランドセーフティ危機に直面した後で導入を行っています。前年度の調査では、同質問に対する回答はそれぞれ13%と8%でした。

処方箋

ブランドに不利益なコンテンツ周辺での露出を回避、画像認識と併用することで相互補完が可能。

考えられる副作用

特になし。

第三者機関による計測

IAS（インテグラル・アド・サイエンス）やDoubleVerify（ダブルベリファイ）といった第三者計測機関をブランドの保護に活用しているマーケターも見られました。予防策として導入しているマーケターは20%で、14%が実際にブランドセーフティ問題に直面した後に導入を行ったと回答しています。前年度の同質問に対する回答はそれぞれ39%と28%で、ここ1年間で大幅な減少傾向にあります。

処方箋

ブランドセーフティに関わる重要な指標を計測・監視。

考えられる副作用

特になし

画像認識は「処方箋」の中で最も将来的に有望なテクノロジーで、前年と比較して活用するブランドが大きく増加しています。現在、マーケターの21%がブランドセーフティ目的で画像認識を導入しており、前年度の12%から大きく伸びました。さらに、16%のマーケターがブランド毀損を実際に体験した後に画像認識を取り入れるようになり、前年度の10%から増加しました。

しかし、画像認識を活用しているマーケターは未だに少数派です。制約の多いブラックリストやホワイトリストへの依存を改善できる有効な手段があるにもかかわらず、その認知や理解はまだ限定的であるといえます。

現在の画像認識技術は、極めて高い精度で物事を正確に分析・特定することができます。アダルト・暴力・競合他社のロゴなど、ブランド毀損リスクの高いビジュアル要素も正確に特定します。画像認識は、ブランド毀損リスクの高い文字情報を検知する自然言語処理と組み合わせると、さらにその効果を発揮します。「ただし、画像認識が唯一絶対というわけではありません」とテクノロジープロバイダーGumGumのコンピュータビジョン責任者であるCambron Carter氏は説明しています。「1つの手法に頼るのではなく、複数の手段を組み合わせることが重要なのです。そうすれば、仮に一つが検知に失敗してももう一方でそれを捕捉することができるのです。」

例えば、航空会社は「墜落」や「テロ」といったフレーズを含むニュースの周辺での広告出稿はなんとかしてでも避けたいと考えます。また、飛行機の残骸の写真が掲載されたニュースの隣に自社広告が掲載されることも必然的に避けたいでしょう。しかし、ニュース記事が画像もしくは文字のどちらか一方にのみ、ブランド毀損リスクの高い内容を含んでいる場合も考えられます。自然言語処理による文脈検知と画像認識を併用することで、こうしたケースにも対応でき、包括的にカバーすることが可能となります。

画像認識と同様、自然言語処理による文脈検知を活用するマーケターも増加しています。21%がブランドセーフティの予防策として導入しており、16%がブランドセーフティ問題が実際に発生した後に導入したと回答しています。前年度の調査では、この割合はそれぞれ13%と8%でした。



FOR Media Buyers

DUE TO Platform & partners

PRIMAR Brand RX

¥

革新的な対処方法

制約の多いブラックリストやホワイトリストが、結果的に望ましいオーディエンスからブランドを引き離してしまっていることは事実ですが、一方で、即効性のあるブランドセーフティ対策としての効果は誰もが認めるところです。「市場はこうした対策も必要としており、マーケターは自分たちがブランドをコントロールするために、とにかくできる手段を講じているのだと思います」とVox MediaのRyan Pauley氏は述べています。ここでの疑問は、望ましいターゲットに適切なリーチ・訴求を行いつつ、同時にブランドを毀損リスクから守るという、針の穴に糸を通すような解決策が見いだせるのかということです。どうやら、答えはテクノロジーにありそうです。

革新的な対処方法

コンピュータビジョン技術や自然言語処理による文脈検知により、いわゆる「コンテクスチュアル（文脈）ターゲティング」機能全般が強化されており、ブラックリストでサイト全体をブロックせずとも、不利益な露出からブランドを守ることができるようになっています。「コンピュータビジョン技術を使うと、画像や動画の中にヌードが含まれていないか、コンテンツの隅の方にナチスのシンボルが含まれていないか、といったあらゆることを調べることができます」とテクノロジープロバイダーGumGumのコンピュータビジョン責任者であるCambron Carter氏は説明しています。コンピュータビジョンによる画像解析と自然言語処理による文脈検知を併用することで、サイト全体を避けずとも、サイト内の特定の部分にブランド露出を限定したり、逆に特定の部分だけにブランドが出ないようにしたり、といったコントロールが可能になります。

SIGNED

BRX

画像認識技術や自然言語処理といった手法は、ブランドにとって好ましいコンテンツを識別することも可能です。例えば、こうした技術を使い、高級車ブランドの広告をコアな自動車ファンの集まるサイトに誘致することもできるのです。

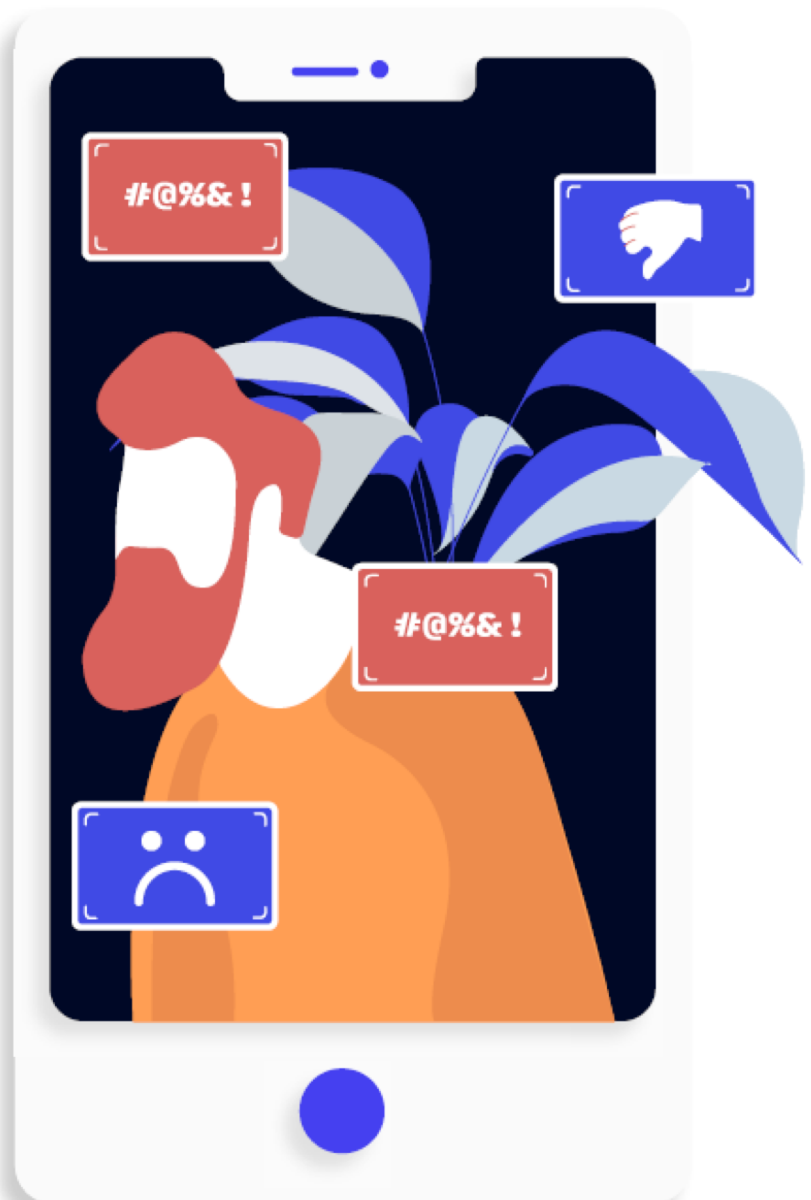
現在の画像認識ツールは主要なDSPに統合されており、何千というパブリッシャーのサイトをリアルタイムにくまなくチェックすることができます。ブラックリストやホワイトリストがある意味一時凌ぎの緩和ケアであるとするれば、画像認識は完治を目指したものであると言えるでしょう。画像認識を利用することで、一見安全そうでないパブリッシャーサイトへも恐れずに思い切って展開することができるのです。

画像認識や自然言語処理といったツールを活用しているマーケターの数は一時的に増加しているわけではありませんが、多くはブランドセーフティを保つ上でこうしたツールが効果的であると認めています。ブランドセーフティにかかる懸念を払拭するのに最も効果的なツールとして、自然言語処理による文脈検知を挙げたマーケターは35%、画像認識を挙げたマーケターは31%に上りました。1年前の調査では、画像認識は18%、自然言語処理による文脈検知は13%という結果でした。

この結果は、一見すると矛盾しているように見えます。こうしたツールが効果的だと回答したマーケターの増加が、実際にツールを採用したマーケターの割合を上回っているのはなぜなのでしょう？ 答えはたった一つです。1年前の調査時点で既にこれらのツールを使用していたマーケターが、今さらながらその価値を実感しているということです。

では、マーケターの多くがまだこうしたテクノロジーを活用していないことはどのように説明できるのでしょうか？ それは単純にこうした新たな技術を知らないということかもしれませんが、予算的な懸念があるからかもしれません。いずれにせよ、機会を逸していることに変わりありません。

「ブランドセーフなコンテンツの意味解析、さらには『センチメント分析』と呼ばれる分析さえも超えたところに私たちは進んでいく必要があります」とBarone氏は指摘しています。画像認識を導入したブランドの多くは、これまでに非常に良好な結果を得ています。



なかなか消えない 症状と新たな 健康不安

広告業界全体としては、まだ完全にブランドセーフティ危機から回復したわけではありません。マーケターの多くは未だに、ここ1年間の間で自社マーケティングコンテンツの周辺に、ヘイトスピーチや暴力やアダルトといったブランドの安全性を損なうコンテンツがあるのを目にすることがあった、と話しています。こうした現状を受け、たとえオーディエンスへのリーチやターゲティングを限定することになってしまっても、制約の多いツールを使ってブランドを守ろうとしているのです。

さらに、ブランドの安全性に対する新たな脅威も出現しつつあります。頭の痛い例として「ディープフェイク」と呼ばれるものが挙げられます。これはAI技術を使い、まるで本物のように見せかけるフェイク動画コンテンツのことを指します。この技術は既に実際に利用されており、ポルノ女優の身体に著名人の顔を合成する、といった利用が行われています。

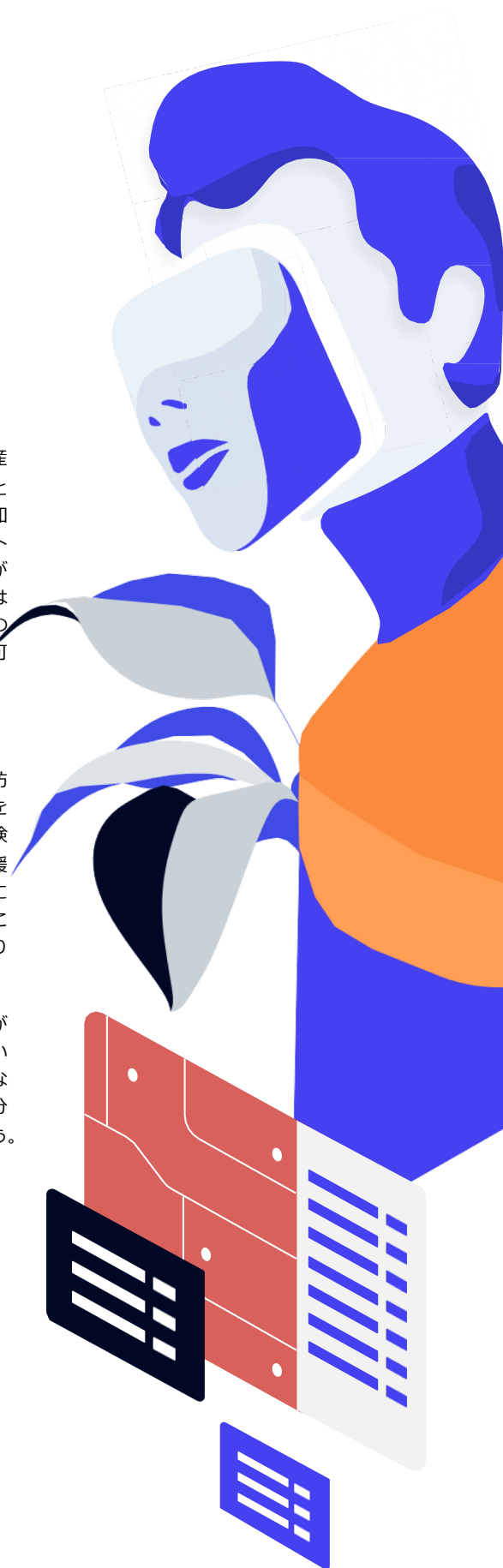
大学のコンピュータサイエンス学部の中には、ディープフェイク動画を自動的に検知できるアルゴリズムを既に開発しているところもあります。ニューヨーク州立大学オルバニ校でコンピュータビジョン・機械学習ラボのディレクターを務めるSiwei Lyu氏によると、ディープフェイクを検知するのに有効なアルゴリズムは既に開発されており、ディープフェイク動画の約99%を検知することが可能とのことでした。

マーケターが心配すべきは、残り1%のディープフェイク動画です。「ディープフェイク動画を探し出し、フレームごとに一つずつ小さな問題を修正している人がいるのを見たことがあります」とLyu氏は述べています。高度なプログラマーが適切なコンピューターを手に入れば、1-2週間のうちにそうした動画を開発することができる、とLyu氏は想定しています。「そうして作られた画像はアルゴリズムで検知するのが難しくなるでしょう」とLyu氏は指摘しています。

今のところ、ディープフェイク動画は好奇心の産物の域を出ていません。GoogleやFacebookといった大企業は効果的にこのような動画を検知し削除しています（両社には、アダルトサイトなどのリサーチを網羅的に行えるだけの体力があります）。しかし、ディープフェイク動画は洗練されてきており、大手プラットフォームの監視の目を掻い潜ることができるようになる可能性は残されています。

米国防高等研究計画局（技術研究を行う米国防総省の機関、DARPA）では既に6,800万ドルを費やして、より洗練されたディープフェイク検知ツールの開発に向けた研究プログラムの支援を行っています。このプログラムは2020年には完了する予定で、その時点で様々な企業がこのツールを商業目的で使用し始めることもあり得る、とLyu氏は指摘しています。

「フェイク動画検知には確実に潜在的な市場が存在していると思います」とLyu氏は述べています。フェイク動画がブランドに及ぼす多大な影響を考えると、広告技術がその市場の大部分を占めるようになることも考えられるでしょう。



完全回復までの 道のり

前年度では、ブランドセーフティ問題はまるで「伝染病」のようなものでした。現在、関係者全員の努力によりこの「伝染病」は収束に向かいつつあります。ソーシャルプラットフォームは多大な努力を行ってより優れた監視や対応を可能にし、論争的になるようなアカウントは閉鎖しました。一方でブランドは、何か問題が起きる前に危険な露出から身を守るべく開発されたツールに注力するようになりました。

全般として市場はいまだに慎重な姿勢を見せていますが、比較的楽観的になってきてはいます。これまでの対応や経過を前向きに受け止めているからです。継続的に免疫システムを強化するために、ブランドセーフティの専門家の役割を強化している企業も多数存在しています。

しかしマーケターの多くは、合理的な予防策と過剰な対応のバランスをどこで合わせるべきか判断しかねています。その結果、もしかしたら取り込めたかもしれない見込み顧客を取り逃がしてしまうということが起きています。こうした機会の損失は、前年度から続いてきた受け入れ難い副作用であり、マーケターが過剰な対策を続ける限り今後も課題として残っていくでしょう。

打開策は最適な技術の活用・応用にあります。画像認識や自然言語処理による文脈検知といった新しいツールや手法に対する理解を深めることで、問題の一部に場当たりに対処するのではなく、総合的な治療・予防を行うことができるようになるでしょう。マーケターには、危険な文脈を避けて身を守ることと、潜在的な顧客層全体に訴求することの両方を実現できる方法が必要です。

その道のりはこれからですが、解決の方向性は既に見えていると言えるでしょう。





gumgum[■]



GumGumについて

@GumGum; gumgum.com

DIGIDAY | CUSTOM

©DIGIDAYメディア 2018年