

Deepfakes and Democratic Integrity: A Policy Framework for AI-Generated Political Content

Christine Corry

University of California: Berkeley

Executive Summary

AI-generated synthetic media has lowered the barrier to producing convincing political disinformation to dollars and minutes. Commercial deepfake-as-a-service tools give low-sophistication actors the ability to mimic primary sources at scale. Human detection of AI-generated multi-modal content performs near chance. AI detection is structurally disadvantaged against iteratively improving generation techniques.

Existing regulatory approaches have not kept pace. Section 230 of the Communications Decency Act preempts state-level content mandates, as California's AB 2655 and AB 2839 demonstrated within weeks of passing. Content-based regulation broadly faces compounding legal fragility from both Section 230 and First Amendment precedent. Voluntary industry commitments signed in 2024 lack enforcement mechanisms, and watermarking and labeling remain broadly unimplemented.

This memo proposes a three-tier content authentication policy: (1) a NIST-administered open watermarking standard, (2) mandatory embedding of that standard alongside C2PA provenance metadata in all commercial AI outputs at the point of generation, and (3) platform-level requirements to surface that data to end users. Enforcement would operate through existing FTC unfair and deceptive practices authority, with third-party technical audits and civil penalties for non-compliance. This approach does not regulate speech or require platforms to moderate content. Instead, it mandates transparency at the API level, where regulatory leverage is highest and legal durability is strongest.

I. Introduction

The modern media landscape is fragmented and decentralized. Social media has cemented itself as a crucial pillar of sources of news in the United States. Pew Research Center found around half (53%) of U.S. adults get news from social media at least sometimes, if not often.¹ The exponential improvement of artificially generated synthetic media has been seeping into the digital landscape and the levers of democracy which are increasingly influenced by it.

In 2024, thousands of people in New Hampshire received a call from the voice of Joe Biden urging Democrats to not vote in the primary election.² Only, former President Biden recorded no such call. Disinformation in political media generally is not novel. Even the election of 1800 used news infrastructure to spread disinformation for political leverage. However, the scale and quality with which realistic disinformation can be created and distributed is new. Convincingly mimicking the likeness of a politician is a fundamentally different threat model for widespread disinformation and deteriorating trust in media. NBC News reported that this particular audio took only 20 minutes and around \$1 to produce.³ Deepfake-as-a-service has decentralized the ability to produce deepfakes giving lower sophistication threat actors the most to gain by producing high-quality false content at scale.

¹ Pew Research Center, "Social Media and News Fact Sheet," 2024.

² Bond, "A Political Consultant Faces Charges," NPR, 2024.

³ Seitz-Wald and Egwuonwu, "A Magician Says a Democratic Operative," NBC News, 2024.

With more of the public receiving news from non-verified or credentialed sources and the barrier to entry to a difficult to recognize deepfake lower than ever, the future of trust and democracy are at risk.

II. Background

The term deepfake refers to digital forgeries which realistically mimic a person's likeness. This policy is primarily concerned with AI-generated deepfakes because of the ease of access to state-of-the-art tools for creating convincing false media at scale.

What is most concerning about deepfake disinformation is the documented inability of humans to recognize AI-generated content. In lab settings, participants only correctly identify an AI-generated voice around 60% of the time.⁴ This is barely better than chance. This study used ElevenLabs, the same company used to generate the Biden calls in New Hampshire. These findings are also in ideal lab settings where participants are specifically asked to evaluate whether or not the voice is AI-generated. This removes the confounding factors of media in which viewers are not primed to consider whether or not the content is real and content may be emotionally charged.

Another study on human accuracy at detecting deepfakes in multi-modal political speech – audio, video, and transcripts – found roughly 51% accuracy on video.⁵ Improved media literacy is not enough to mediate this risk, nor is simply augmenting human detection abilities with AI detection tools.⁴ Policy targeting the generation, posting, or hosting of political deepfake disinformation must be enacted.

III. Assessment of Current Approaches

A. First amendment protections

A prominent example cited as a type of protected free speech which could be encroached upon by regulation of political disinformation deepfakes at the creation and posting level is satire. This nuanced balance of regulating lies and protecting harmless speech has played out recently with the Stolen Valor Act, which originally made it a crime to falsely claim military medals in any context. In 2012, the Supreme Court struck down the Stolen Valor Act claiming violation of the First Amendment in *United States v. Alvarez*. The majority opinion states that First Amendment protections extend to even knowingly false speech unless legally recognized harm is caused. They found the wording of the Stolen Valor Act to be too broad, threatening to suppress trivial lies.⁶ In 2013, a revised version of the act was signed, criminalizing lies about military service with intent to obtain money, property, and other tangible benefits.

The basis of the Supreme Court ruling supposes a free marketplace of ideas in which falsities can be self-corrected. For claims of receiving the Medal of Honor, there is a well documented source of truth. One can search online for the real list of honorees. Deepfakes are much more difficult to disprove and distinguish from ground truth images. They are not subject to the same forces of a free marketplace because of the speed and decentralization of the platform media environment and difficulty to detect and disprove.

The revised act shows the government is open to regulation on intentional disinformation resulting in demonstrable harm. The same framework could be extended to AI-generated political disinformation, although recent moves by the Federal Election Commission (FEC) dampen this possibility considerably.

⁴ Barrington, Cooper, and Farid, "People Are Poorly Equipped," *Scientific Reports*, 2025.

⁵ Groh et al., "Human Detection of Political Speech Deepfakes," *Nature Communications*, 2024.

⁶ FIRE, "United States v. Alvarez."

In 2024, Governor Ron DeSantis of Florida used deepfakes in a campaign against President Donald J. Trump.⁷ An immediate response from the FEC voted to evaluate rules concerning AI use in campaigns. However, they later decided no new rules were needed since existing technologically neutral regulations on fraudulent misrepresentation would sufficiently cover AI-generated content. The existing approach does not account for the evidentiary problem of proving a specific piece of content is AI-generated with authentication infrastructure.

B. Legal liability blocked by Section 230

The primary policy roadblock for effective deepfake regulation at the level of dissemination is Section 230 of the Communications Decency Act. This 1996 regulation bars liability for providers of content posted to their platforms. This protection made sense in the era of a limited scope internet in which content providers and platforms were small and few, based on the principle that platforms do not engage in editorial review over content. In the modern internet, platforms are powerful and pervasive. Unlike a newspaper editor who bears liability for published content, or a newsstand owner who faces consequences once notified of defamatory material, a platform under Section 230 bears no legal consequence regardless of what it hosts or how aggressively its algorithm distributes particular content, thus has no liability incentive to moderate.

Some scholars argue that implementation of algorithmic amplification on social media platforms makes for a kind of editorial judgement which carries publisher-like liability. This is still an open question in U.S. law. The Supreme Court had the opportunity to decide whether Section 230's liability shield covered algorithmic recommendations of user content in 2023 in *Gonzalez v. Google LLC*.⁸ Instead, they avoided Section 230 altogether, ruling instead on liability under the Anti-Terrorism Act.

More recent strides towards platform litigation have been made in March of 2026 with *K.G.M v. Meta & Google*. The California Superior Court found Meta and Google negligent and ruled that Instagram and YouTube were deliberately built to be addictive.⁹ While this move does not open the floodgates for broader regulation since it focuses on conduct and design rather than content, it does give momentum in a regulatory direction and reveal vulnerabilities for platform liability. This case proves the legal viability of the conduct/content distinction as a relevant policy lever.

C. State-level content regulation

Many states have attempted regulation in this problem space. California is particularly well suited to lead this charge since it contains the Silicon Valley technology hub. AB 2655, the Defending Democracy from Deepfake Deception Act, would have required large online platforms to “block the posting of materially deceptive content related to elections in California, during specified periods before and after an election” and label certain additional content as fake during the same periods.¹⁰ AB 2839 would have targeted the content distributors more broadly, prohibiting “a person, committee, or other entity from knowingly distributing an advertisement or other election communication, as defined, that contains certain materially deceptive content.”¹¹

⁷ Wines, "DeSantis Campaign Uses Deepfake," *New York Times*, 2023.

⁸ *Gonzalez v. Google LLC*, 598 U.S. 617 (2023); Feder and Gillespie, "The U.S. Supreme Court Punts on Section 230," Covington, 2023.

⁹ Contreras, "Meta, YouTube Found Liable," NPR, 2026.

¹⁰ California A.B. 2655, 2023–2024 Reg. Sess. (Cal. 2024).

¹¹ California A.B. 2839, 2023–2024 Reg. Sess. (Cal. 2024).

Enforcement of both CA AB 2839 and, soon after, 2655 was blocked by a U.S. district judge, Senior Judge John Mendez, only weeks after it was signed by Governor Newsom.¹² The challenges to the bills were brought forward with concerns over free speech, but Mendez steered mostly clear of that issue, instead highlighting the federal preemption of Section 230.¹³ The First Amendment considerations are a deeper constitutional issue, while Section 230 is the more immediate procedural block. It is unclear what impact First Amendment precedent would have on similar legislation which does not violate Section 230.

Federal preemption is also a threat for state-level regulation. The House-passed version of the Big Beautiful Bill included a moratorium on state-level AI regulation.¹⁴ The Senate voted almost unanimously to strip the moratorium.¹⁵ These actions demonstrate both the federal threat to a patchwork of state regulation and bipartisan interest in AI regulation, at least in the Senate.

D. Voluntary industry standards

In 2024, tech giants Microsoft, Meta, Google, Amazon, X, OpenAI, and Tiktok voluntarily agreed to an industry “accord” to mediate the risk of AI-generated political disinformation. Importantly, it included commitments to develop watermark technology and the detection and labeling of realistic AI-generated content.¹⁶ The agreement lacks accountability mechanisms and adequate watermarking and deep fake labeling are broadly not followed.¹⁷ ¹⁸ Industry voluntary agreements are not sufficient; disinformation mediation standards must be well defined and have accountability mechanisms with third-party confirmation of compliance.

E. Detection as a strategy

The reactive strategy of deepfake detection is not sustainable given the speed with which this technology develops. Detection performance can drop up to 50% when tested against real-world deepfakes rather than academic datasets.¹⁹ As soon as detection services are deployed at scale, they are used as the benchmarks for improved generation technology. There is a structural advantage in the generation algorithm iteratively improving against new detection methods while detection based on novel analytic methods lags behind.

F. AI-Generated content regulation trends

Other recent legislation signals a different current for policy paths towards regulating particular kinds of AI-generated media. The TAKE IT DOWN Act criminalized publication of deepfakes depicting intimate visual content, generally sexually explicit in nature. Victims are empowered to demand removal of such content. Platforms could face FTC enforcement for failing to remove offending content. The DEFIANCE Act creates a “private right of action for the production, possession, solicitation, or disclosure of unconsented AI deepfakes that depict intimate visual content.”²⁰

Congress has found a path for creating enforceable individual rights around intimate imagery deepfake

¹² Ballotpedia, "AI Deepfake Policy in California."

¹³ Gurman, "Musk's X Wins Court Challenge," *Politico*, 2025.

¹⁴ Goodwin Law, "House Passes 10-Year Federal Moratorium," 2025.

¹⁵ Reints, "Senators Reject 10-Year Ban," *Time*, 2025.

¹⁶ Treisman, "Tech Giants Lay Out Plan," NPR, 2024.

¹⁷ Brennan Center, "New Tech Accord to Fight AI Threats," 2024.

¹⁸ Rijsbosch, van Dijck, and Kollnig, "Missing the Mark," *Policy & Internet*, 2026.

¹⁹ Chandra et al., "Deepfake-Eval-2024," arXiv, 2025.

²⁰ Brown Rudnick, "Legislating Deepfakes and AI Tort Liability," 2025.

content where First Amendment tensions are less acute and bipartisan political will is stronger compared to political disinformation deepfakes.

IV. Proposed Policy

The First amendment protections and liability shield of Section 230 make content-based regulation legally fragile. Mandates on firms directly are much more feasible, especially in the climate of increasing bipartisan awareness of AI risks broadly.

I propose a three-tier content authentication stack: (1) a NIST-administered open, interoperable watermarking standard, (2) mandatory embedding of both that standard and C2PA provenance metadata in all outputs from commercial AI systems at the point of generation, and (3) platform-level verification requirements to surface provenance data and watermarking results to end users. This is a project in transparency more than content specific regulation.

C2PA, the open provenance standard, carries forward the data's history for digital media files including where they originated from and how they were edited. The information is cryptographically embedded in the file. There are other potential authenticity approaches, but C2PA is cited in the voluntary agreement signed by leading cross-industry bodies and a free, open, and accessible standard.

There is not currently an open standard for watermarking AI generated content. A proprietary watermarking system which the open standard could mirror is Google's SynthID. SynthID is an invisible watermark specifically designed to be resistant to metadata stripping transformations. It is demonstrably robust, achieving up to 91% detection confidence after an extended chain of transformations: screenshotting, cropping, compressing, and re-uploading to Instagram. The combination of techniques gives a robust multi-use-case approach. Compelling SynthID to be made open, free, and available across the industry would be difficult, motivating the creation of an open standard for the purpose of interoperable use. NIST's development of an open standard for watermarking would be enforced by Congressional mandate with a defined timeline and conformance certification program.

The enforcement point for mandatory C2PA and watermarking is the API call to the closed model. Each generation passes through the firm's infrastructure, creating a bottleneck at which these regulations must be enforced. Existing FTC unfair and deceptive practices authority would be sufficient to require periodic third-party technical audits to confirm generated content carries C2PA and watermarking data with civil penalties per violation and market access conditionality for repeated non-compliance.

This provenance data and watermarking is not directly accessible or verifiable by individual users which requires platforms to include verification software. The content itself is not being regulated, so First Amendment speech would not be suppressed and platforms are not made to mediate content, only surface the provenance data. Periodic conformance certification would handle the baseline technical implementation verification and would be conducted under FTC unfair and deceptive practices authority. Independent red-team auditing could reinforce this accountability measure in the wild. Particularly intense auditing should take place around election cycles to ensure deepfake political disinformation is not going undetected.

Similar legislation is passed but not yet enforced in California's AI Transparency Act which was signed into law but whose operative date was pushed back to August 2026 in alignment with the EU AI Act. This requires providers to provide a free, publicly accessible AI detection tool, optional visible label for AI content, and mandatory hidden watermark. These requirements are firm specific rather than cross-platform with no broad standards on the interoperability or transparency of the watermarking

technique and detection tool. This fragmentation divides resources for technological development to make watermarking and provenance data attachments resistant to higher sophistication threat actors. It also makes mandatory verification on the part of platforms less straightforward with multiple software implementations rather than two industry standards.

V. Potential Barriers

A. Political and Industry Opposition

Industry resistance is the most predictable barrier. Firms running high-volume APIs may argue that watermarking adds computational overhead at inference time and that C2PA manifest signing requires maintaining certificate infrastructure, both of which impose latency and compliance costs. Smaller AI firms and startups could argue that the compliance burden is proportionally heavier for them than large firms, like Google and OpenAI, who have already invested in provenance data and watermarking infrastructure. However, compliance infrastructure is already being built into major commercial API providers in response to voluntary commitments and California's AI Transparency Act. There is also a potential argument to be made about competitive advantage between US-regulated firms and foreign competitors not covered by these requirements. This applies equally to all AI regulation and argues for international coordination rather than against domestic action.

AI firms have a particular incentive to keep standards practices at the industry level rather than ceding them to the government. OpenAI spent \$1.76 million on AI policy government lobbying in 2024, almost seven times as much as was spent on lobbying in 2023.²¹ Meta's Super PAC, American Technology Excellence Project, supports tech-friendly candidates in state elections and opposes emerging state AI regulation. OpenAI president, Greg Brockman, and venture capital firm a16z invested \$100 million in Leading the Future, a PAC advocating against strict AI regulation.²² Large tech firms including Google, Meta, Amazon, and Microsoft lobbied in favor of the House package which called for a 10-year ban on state AI regulation. Firms have successfully recast mandatory standards as threats to national competitiveness rather than consumer protections. Voluntary commitments without accountability mechanisms are an instrument of this reframing – absorbing political pressure for action without altering underlying practices or ceding agenda-setting power to regulators. A mandatory federal authentication standard would represent exactly the kind of government agenda-setting these firms have spent hundreds of millions of dollars to prevent.

A similar dedication to industry innovation positioned against regulation is held by the Trump administration. In a December 2025 executive order, the Trump administration stated plainly “to win, United States AI companies must be free to innovate without cumbersome regulation.”²³ The Biden-era executive order directing NIST developments on guidelines for generative AI and synthetic content was rescinded on the first day of the Trump administration.²⁴ Federally, the U.S. is bound up in race dynamics and believes that guardrails will only slow development and advantage foreign AI developers. Even if this regulation was to pass through Congress, surviving Presidential veto may be a challenge.

B. Structural Legislative Barriers

The seasonality of attention towards this issue is also a structural barrier. Political disinformation and its effects are the most prominent before elections. Urgency fades after elections and pushes toward

²¹ Hao, "OpenAI Ups Its Lobbying Efforts Nearly Seven-Fold," MIT Technology Review, 2025.

²² Gold, "Meta Launches Super PAC to Fight AI Regulation," Axios, 2025.

²³ Trump, "Eliminating State Law Obstruction of National Artificial Intelligence Policy," White House, 2025.

²⁴ NIST, "Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence" (*rescinded January 20, 2025*).

substantive legislation are watered down or stalled. The mismatch between political salience and legislative timelines is an implicit barrier to passing this regulation federally. Even legislation introduced today would face a multi-year implementation timeline; NIST standardization alone would take years. The 2026 midterm cycle would pass before any mandated standard is created, no matter the actions of Congress.

C. Technical Scope Limitations

A potential limitation of regulation at the level of AI firms is that it does not cover open-weight models run locally, but deepfakes-as-a-service for low sophistication threat actions would be impacted. Most examples of widespread political disinformation deepfakes used commercial APIs; Biden's voice clone audio was created at ElevenLabs. Generally, locally-run open-weight models are lower quality, technically demanding to run at scale, and more accessible to sophisticated threat actors. Additionally, the absence of provenance metadata can itself be a signal for scrutiny. This sort of positive signal for authentic content with the C2PA metadata inverts the liar's dividend. Authentic content cannot be falsely claimed as AI-generated and media without authentication metadata will face a higher evidentiary burden.

U.S. based regulation and enforcement at the firm level also does not cover foreign models. For example, a person in the U.S. using the DeepSeek API to generate synthetic political media is outside the regulatory perimeter. The same argument for the inverted liar's dividend is also relevant to the domestic misuse threat model, however foreign state-sponsored operations are a different sort of threat entirely and not covered by this regulation. Bilateral AI governance agreements are better suited to this model.

VI. Conclusion

AI-generated political disinformation is a threat to democracy and trust in media broadly. Regulation of content is legally fragile with First Amendment protections and Section 230's liability shield over platforms. California's attempts at this issue have been blocked by such tensions and are threatened by federal preemption. Voluntary industry standards are demonstrably insufficient without accountability measures and detection is a perpetual game of cat-and-mouse. A three-tier policy approach establishing open standards for provenance metadata and watermarking attachment and verification with enforced use by both AI-firms generating content and platforms which may host it circumvents content-based legal constraints and creates an inverted evidentiary burden – the absence of authentication data is cause for scrutiny – using existing FTC authority for enforcement.

Sources

Ballotpedia. "AI Deepfake Policy in California." Accessed May 2026.

https://ballotpedia.org/AI_deepfake_policy_in_California.

Barrington, S., E.A. Cooper, and H. Farid. "People Are Poorly Equipped to Detect AI-Powered Voice Clones." *Scientific Reports* 15, no. 11004 (2025). <https://doi.org/10.1038/s41598-025-94170-3>.

Bond, Shannon. "A Political Consultant Faces Charges and Fines for Biden Deepfake Robocalls." NPR, May 23, 2024.

<https://www.npr.org/2024/05/23/nx-s1-4977582/fcc-ai-deepfake-robocall-biden-new-hampshire-political-operative>.

Brennan Center for Justice. "New Tech Accord to Fight AI Threats in 2024 Lacks Accountability." February 2024.

<https://www.brennancenter.org/our-work/analysis-opinion/new-tech-accord-fight-ai-threats-2024-lacks-accountability-companies>.

Brown Rudnick. "Legislating Deepfakes and AI Tort Liability." October 2025.

<https://briefings.brownrudnick.com/post/102lqka/legislating-deepfakes-and-ai-tort-liability>.

California Assembly Bill 2655, Defending Democracy from Deepfake Deception Act, 2023–2024 Regular Session. Cal. 2024.

https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202320240AB2655.

California Assembly Bill 2839, 2023–2024 Regular Session. Cal. 2024.

https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202320240AB2839.

Chandra, Nuri, et al. "Deepfake-Eval-2024: A Multi-Modal In-the-Wild Benchmark of Deepfakes Circulated in 2024." arXiv:2503.02857, March 2025. <https://arxiv.org/abs/2503.02857>.

Contreras, Russell. "Meta and YouTube Found Liable in Landmark Social Media Addiction Trial." NPR, March 25, 2026. <https://www.npr.org/2026/03/25/nx-s1-5746125/meta-youtube-social-media-trial-verdict>.

Feder, Janis, and Tom Gillespie. "The U.S. Supreme Court Punts on Section 230 in *Gonzalez v. Google LLC*." Covington & Burling, May 2023.

<https://www.cov.com/en/news-and-insights/insights/2023/05/the-us-supreme-court-punts-on-section-230-in-gonzalez-v-google-llc>.

Foundation for Individual Rights and Expression (FIRE). "United States v. Alvarez." Accessed May 2026.

<https://www.fire.org/supreme-court/united-states-v-alvarez>.

Gold, Ashley. "Meta Launches Super PAC to Fight AI Regulation as State Policies Mount." *Axios*, September 23, 2025. <https://www.axios.com/2025/09/23/meta-superpac-ai-regulation>.

Gonzalez v. Google LLC, 598 U.S. 617 (2023). <https://www.law.cornell.edu/supct/cert/21-1333>.

Goodwin Law. "House Passes 10-Year Federal Moratorium on State AI Regulation." May 2025.

<https://www.goodwinlaw.com/en/insights/publications/2025/05/alerts-practices-aiml-house-passes-10-yea>

[r-federal-moratorium](#).

Groh, M., A. Sankaranarayanan, N. Singh, D.Y. Kim, A. Lippman, and R. Picard. "Human Detection of Political Speech Deepfakes Across Transcripts, Audio, and Video." *Nature Communications* 15, no. 7629 (2024). <https://doi.org/10.1038/s41467-024-51998-z>.

Gurman, Mark. "Musk's X Wins Court Challenge to California Deepfake Law." *Politico*, August 5, 2025. <https://www.politico.com/news/2025/08/05/elon-musk-x-court-win-california-deepfake-law-00494936>.

Hao, Karen. "OpenAI Has Upped Its Lobbying Efforts Nearly Seven-Fold." *MIT Technology Review*, January 21, 2025. <https://www.technologyreview.com/2025/01/21/1110260/openai-ups-its-lobbying-efforts-nearly-seven-fold/>.

National Institute of Standards and Technology. "Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence." Accessed May 2026. <https://www.nist.gov/artificial-intelligence/executive-order-safe-secure-and-trustworthy-artificial-intelligence>.

Pew Research Center. "Social Media and News Fact Sheet." Pew Research Center Journalism Project, 2024. <https://www.pewresearch.org/journalism/fact-sheet/social-media-and-news-fact-sheet/>.

Reints, Renae. "Senators Reject 10-Year Ban on State-Level AI Regulation." *Time*, July 1, 2025. <https://time.com/7299044/senators-reject-10-year-ban-on-state-level-ai-regulation-in-blow-to-big-tech/>.

Rijsbosch, Bram, Gijs van Dijck, and Konrad Kollnig. "Missing the Mark: Adoption of Watermarking for Generative AI Systems in Practice and Implications Under the New EU AI Act." *Policy & Internet* (April 2026). <https://doi.org/10.1002/poi3.70041>.

Seitz-Wald, Alex, and Nnamdi Egwuonwu. "A Magician Says a Democratic Operative Paid Him to Make the Fake Biden New Hampshire Robocall That Is Under Investigation." *NBC News*, February 24, 2024. <https://www.nbcnews.com/politics/2024-election/biden-robocall-new-hampshire-strategist-rcna139760>.

Treisman, Rachel. "Tech Giants Lay Out Plan to Fight AI Election Deepfakes." *NPR*, February 16, 2024. <https://www.npr.org/2024/02/16/1232001889/ai-deepfakes-election-techaccord>.

Trump, Donald J. "Eliminating State Law Obstruction of National Artificial Intelligence Policy." Presidential Action, White House, December 11, 2025. <https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

Wines, Michael. "DeSantis Campaign Uses Deepfake Images of Trump and Fauci in Attack Ad." *New York Times*, June 8, 2023. <https://www.nytimes.com/2023/06/08/us/politics/desantis-deepfakes-trump-fauci.html>.