

YOU (ai)N'T HEARD NOTHIN' YET

A NEW FRAMEWORK FOR FORENSIC AUDIO AUTHENTICATION IN THE AGE OF AI

I. INTRODUCTION: THE FORENSIC CRISIS OF AUTHENTICITY

Forensic science exists to serve a singular purpose: to provide courts with reliable, factual evidence that assists in the search for truth. That mission now faces a behemoth. Artificial intelligence can generate convincing voice clones from mere seconds of audio,¹ while courts still admit voice identification based on lay familiarity, a standard conceived for an analog era.²

Empirical studies have detailed this divide, citing human listeners detect sophisticated audio deepfakes with only about 73% accuracy,³ an error rate inconsistent with modern forensic practice. By contrast, state-of-the-art forensic systems often reach 85–95% accuracy in controlled settings,⁴ yet they face intense gatekeeping scrutiny under *Kelly/Frye* (if novel) and *Sargon*,⁵ and are not routinely deployed for ordinary cases. This inversion – admitting what is least reliable while burdening what is more reliable – creates a gap that enables two catastrophic failures.

In *Mendones v. Cushman & Wakefield*, parties submitted AI-generated audio testimonials that passed authentication and reached summary judgment before being exposed, resulting in terminating sanctions.⁶ Conversely, in *Huang v. Tesla, Inc.*, Tesla invoked the *possibility* of deepfakes to challenge a 2016 recording of CEO Elon Musk, despite the video predating the term “deepfake” and being witnessed live by journalists.⁷ The judge rejected the gambit as “deeply troubling,” noting that Tesla’s position would allow public figures to “avoid taking ownership of what they did actually say and do.”⁸

Together, these cases illustrate the crisis: courts lack structured tools to detect fabrications *or* to refute baseless deepfake allegations. The result is Chesney and Citron's "liar's dividend": technology-enabled doubt that undermines trust in all audio evidence, authentic or not.⁹

This article proposes a practical solution: the Graduated Audio Authentication Framework. It calibrates the level of forensic scrutiny to the risk and consequence of error, integrates with existing rules, and remains adaptable as technology evolves.

II. MOTIVATING CASES: THE TWO-SIDED AUTHENTICATION FAILURE

Current authentication doctrine fails in both directions: it admits fabricated audio while enabling baseless challenges to authentic recordings.

A. WHEN FABRICATED AUDIO PASSES: MENDONES V. CUSHMAN & WAKEFIELD

The *Mendones* litigation illustrates the failure to detect. Plaintiff submitted AI-generated testimonials with classic indicators of synthesis, such as robotic prosody and unnatural delivery. Discovery did not surface the problem; it emerged only at summary judgment, prompting terminating sanctions.¹⁰

Framework Application: A pretrial authentication conference would have assigned these exhibits to Tier 3 based on centrality and witness unavailability. The proponent would then need a qualified expert and multi-method analysis. Early, structured scrutiny likely would have resulted in exclusion for lack of authentication, avoiding wasted resources.

B. WHEN AUTHENTIC AUDIO IS CHALLENGED: HUANG V. TESLA, INC.

In *Huang v. Tesla, Inc.*, plaintiffs sought to introduce Elon Musk's public statements promoting Autopilot's safety.¹¹ Tesla opposed, arguing that Musk "is the subject of many 'deepfake' videos" and therefore could not verify the statements' authenticity.¹² This argument

ignored that the recording predated widespread deepfake technology and was witnessed in person by journalists.¹³

Judge Evette D. Pennypacker partially rejected Tesla's argument, writing that "What Tesla is contending is deeply troubling."¹⁴ Excluding the evidence, she ordered a deposition and noted that such a position would allow public figures to "hide behind the potential for their recorded statements being a deep fake to avoid taking ownership of what they did actually say and do."¹⁵

Framework Application: Under the Graduated Framework, Tesla's challenge would trigger tier assignment based on allegation specificity. A vague claim that statements "might be" deepfakes, without forensic analysis or temporal impossibility evidence, would warrant Tier 2 at most: basic technical screens sufficient to confirm authenticity. If Tesla alleged a *specific* deepfake tool or provided forensic evidence, Tier 3 would apply. The Framework converts speculative deepfake allegations into testable evidentiary claims with proportional burdens, preventing weaponized liar's dividend tactics.

III. SCIENTIFIC FOUNDATIONS: WHAT COURTS NEED TO KNOW

Three empirical pillars justify calibrated authentication requirements:

1. **Human perception is unreliable.** Human accuracy hovers around 73% overall,¹⁶ and can fall near chance in harder regimes. Confidence can even inversely correlate with accuracy for familiar voices.¹⁷ To contextualize this: eyewitness misidentification, which contributes to approximately 69% of DNA exonerations stems from similar perceptual overconfidence.¹⁸
2. **Detection models are brittle.** Self-supervised models that excel on known synthesis methods can degrade severely on unknown methods or in the wild,¹⁹ sometimes collapsing toward 50–60% accuracy.²⁰
3. **Neural vocoder artifacts are a strong signal.** Modern text-to-speech (TTS) pipelines leave subtle, characteristic spectral fingerprints that are algorithmically detectable and comparatively robust for screening.²¹

A. GATEKEEPING PRINCIPLES

Courts do not need a signal-processing seminar, but they do need modern reliability concepts to exercise effective gatekeeping.

- **Bench vs. Field Performance:** Published accuracies typically reflect controlled benchmarks.²² Real-world performance drops when conditions shift.²³ Gatekeeping should probe whether the expert validated methods on data resembling the case conditions.
- **Report Error in Context:** Binary “accuracy” can mask tradeoffs. Forensic testimony should report error rates using accepted measures (e.g., Equal Error Rate or EER),²⁴ allowing courts to weigh false-positive versus false-negative risk consistent with probative value balancing.²⁵
- **Method Calibration:** Experts should disclose training/validation splits and whether independent, out-of-domain tests were performed.²⁶

B. THE DEFENDER’S DILEMMA: WHY DETECTION LAGS BEHIND SYNTHESIS

A structural asymmetry defines the forensic landscape: synthesis models optimize for deception, while detectors optimize for discrimination. This “Defender’s Dilemma” explains why perfect detection is mathematically elusive.

- **The Optimization Gap:** Generative adversarial networks (GANs) and diffusion models are explicitly trained to minimize the perceptual difference between real and fake audio.²⁷ As synthesis improves, artifacts are pushed into subtle “latent” spaces that are increasingly difficult to distinguish from natural variability.²⁸
- **The Generalization Problem:** Detection models suffer from a “closed-set” limitation. They learn to spot the specific artifacts of known synthesis tools (e.g., ElevenLabs v1). When a new tool (e.g., ElevenLabs v2, v3, or other proprietary model) emerges with different artifact patterns, the detector’s performance can collapse.²⁹ This is not a failure of the expert; it is a property of the domain.
- **Implication for Courts:** Judicial expectation of “100% certainty” is scientifically illiterate in this context. The goal is not certainty but *calibrated probability*. A conclusion that “this audio contains artifacts consistent with diffusion-based synthesis” is highly probative even if the specific model cannot be identified. Conversely, the absence of artifacts does not prove authenticity, it merely fails to disprove it. This asymmetry necessitates the “multi-method fusion” approach described below.

C. KNOWN FAILURE MODES (AND HOW TO MITIGATE THEM)

Understanding *why* models fail is essential for evaluating expert reliability.

- **Cross-Dataset Collapse:** Detectors trained on a narrow family of synthesis pipelines (e.g., ElevenLabs) often underperform on unknown pipelines (e.g., Tortoise or VALL-E).³⁰

- **Mitigation:** Courts should require cross-domain validation and, when feasible, the use of two conceptually different detectors (e.g., vocoder artifacts plus SSL anomalies) to reduce correlated error.
- **Compression and Channel Effects:** Heavy compression (WhatsApp, Zoom), telephony bands, and aggressive noise suppression can mask or mimic artifacts.
 - **Mitigation:** Experts must analyze the native file, avoid multiple generations, document the codec chain, and use pre-processing matched to the file's actual signal characteristics.
- **Language and Dialect Mismatch:** Accuracy can vary significantly by language and dialect.³¹
 - **Mitigation:** Prefer models trained on or adapted to the spoken variety; interpret results conservatively when mismatched.
- **Adversarial Examples:** Targeted irregularities can fool some detectors.³²
 - **Mitigation:** Ensemble methods and robustness checks (e.g., slight re-encoding) to see if a conclusion is stable.

C. FROM SCREENING TO PROOF: MULTI-METHOD FUSION

The Framework encourages combining methods that fail in different ways. A practical sequence for Tier 3 looks like this:

1. **Intake and Triage:** Document chain of custody, recording date, device type, and codec. These factors inform tiering.
2. **Low-Cost Screens:** Run spectrogram scans and metadata checks; apply a modern vocoder-artifact screen.³³
3. **Deep Analysis:** If risk remains high, add SSL-based detection validated for the case conditions,³⁴ and attempt temporal localization to tie conclusions to the specific sounds.³⁵
4. **Quantify and Visualize:** Report findings with operating thresholds and uncertainty; include anomaly maps and controlled standards.

Experts should explicitly explain how each method's limitations affect the bottom-line opinion. That is what *Sargon* requires: not perfection, but reliable methods properly applied.³⁶

D. CHOOSING MODELS TO MATCH CASE CONDITIONS

Model choice is dispositive. A detector trained on high-fidelity broadcast speech may fail on telephony audio. Experts should inventory case conditions, language, dialect, sample rate, codec, SNR, background noise, device class, and select or adapt models accordingly. Gatekeepers should view "black box" applications of generic models with skepticism. When localization is

outcome-determinative, the expert must confirm that the model's temporal resolution is sufficient for the disputed phrase length.

IV. THE GRADUATED AUDIO AUTHENTICATION FRAMEWORK (GAAF)

The Framework aligns scientific scrutiny with risk and consequence, applying proportionality principles embedded in California's evidence and discovery statutes.³⁷

A. COUNTERING THE LIAR'S DIVIDEND: A POLICY IMPERATIVE

The "liar's dividend", the ability to disclaim authentic evidence by invoking the mere possibility of synthesis, is a central harm of deepfake technology.³⁸ *Huang* demonstrates this dynamic: Tesla avoided ownership of public statements by hiding behind posed deepfake claims.³⁹

If courts accept the *Huang* argument, that the mere existence of deepfake technology renders all recordings inherently suspect, they effectively repeal the Best Evidence Rule for the digital age. The Rule's historic purpose was to prevent fraud by requiring the "original,"⁴⁰ but deepfakes challenge the very ontology of an original. Under the *Huang* regime, no recording could ever be authenticated without the creator's admission, returning the legal system to a pre-modern reliance on fallible human memory. This weaponized skepticism threatens to erase decades of progress in evidentiary reliability.

The GAAF counters this by converting speculative allegations into tiered, testable claims. A party asserting fabrication must either (a) provide a credible evidentiary basis that triggers Tier 3 or 4 authentication, or (b) accept that Tier 1 or 2 authentication suffices. Costs shift to the challenger if basic technical screens confirm genuineness. This prevents the liar's dividend from functioning as cost-free obstruction while restoring trust in authenticated evidence through transparent, auditable findings.

B. THE TIERS

- **Tier 1 (Low Risk):** Lay authentication under standard rules suffices where synthesis is implausible given circumstances and corroboration is strong.⁴¹
- **Tier 2 (Moderate Risk):** Lay testimony must be supplemented with targeted technical checks (spectrogram, metadata, neural vocoder screening) and a short technician report.
- **Tier 3 (High Risk):** Comprehensive, multi-method analysis by an expert qualified under expert witness statutes (spectral analysis, vocoder artifact detection, SSL-based detection), with quantified findings and method limits.
- **Tier 4 (Very High Risk):** All Tier 3 requirements plus independent verification,⁴² or a court-appointed neutral.⁴³

C. TIER ASSIGNMENT FACTORS

Courts should make tier assignments on the record using factors that reflect risk:

- **Centrality:** Does the recording substantially influence liability or liberty? Criminal guilt-phase evidence is presumptively Tier 3 when contested.
- **Allegation Strength:** Is there a direct, credible allegation of synthesis (e.g., identified tool, timing impossibility), or only a speculative concern?
- **Recording Conditions:** Low quality or custody gaps justify higher tiers.
- **Temporal Feasibility:** Post-2020 recordings receive heightened scrutiny if challenged.⁴⁴

D. ILLUSTRATIVE APPLICATIONS

1. Criminal: The Jailhouse Confession

Suppose a criminal defendant alleges a jailhouse confession recording was fabricated using a cloning tool.

- **Tiering:** Central to outcome + specific allegation = **Tier 3**.
- **Methods:** The expert verifies hashes and runs screens. Vocoder screening shows elevated artifacts. SSL-based detection flags spoof likelihood around the disputed sentence “I did it.” Localization places anomalies in an abnormal window.
- **Opinion:** The expert reports that the disputed phrase is likely synthetic, supported by convergent methods with stated error rates.⁴⁵

2. Civil: The Whistleblower vs. The Corporation (Addressing Asymmetry)

In *Doe v. MegaCorp*, a whistleblower introduces a recording of the CEO admitting to fraud. MegaCorp, possessing vast resources, claims the recording is a deepfake and demands exhaustive, costly analysis to drain Doe’s litigation budget.

- **Tiering:** High stakes + “Deep Pockets” asymmetry = **Tier 3**.
- **The Resource Fix:** The court applies the Framework’s cost-shifting leverage. Since the initial Tier 2 screen (performed by Doe’s low-cost expert) shows no obvious synthesis, the court orders that MegaCorp must cover the cost of the Tier 3 independent analysis.⁴⁶
- **Result:** If the recording is authentic, MegaCorp bears the full cost, plus potential sanctions for the frivolous challenge. If it is fake, costs shift back to Doe. This “put up or shut up” mechanism prevents wealthy litigants from using authentication challenges as a weapon of attrition.

V. PROCEDURAL IMPLEMENTATION

1. *PRETRIAL AUTHENTICATION CONFERENCE*

Parties should identify contested audio and propose tier assignments early in the litigation. This prevents late-stage crises and aligns effort with risk.⁴⁷ The conference serves as the initial gatekeeping step, allowing the court to filter out frivolous “liar’s dividend” challenges (assigning them to Tier 1 or 2) while ensuring genuine disputes receive adequate resources (Tier 3 or 4).

2. *TIER-CALIBRATED EXPERT STANDARDS*

The Framework distinguishes between technicians and scientists:

- **Tier 2 Technician:** Competent to run standard commercial screening tools and interpret basic metadata. Qualifications may include certification in digital forensics or specific tool training.
- **Tier 3/4 Scientist:** Must be capable of withstanding rigorous Kelly/Frye and Sargon scrutiny.⁴⁸ This requires more than just tool proficiency; the expert must understand the underlying architecture, be able to explain training data and failure modes, and ideally have peer-reviewed contributions or deep technical experience in audio synthesis detection. They must be able to defend why a specific model is appropriate for the specific case conditions.

3. *ROBUST DISCOVERY AND REPLICATION*

Transparency is the antidote to “black box” forensics. For Tier 3 and 4 cases, discovery must go beyond the final report. To enable meaningful cross-examination and replication, proponents must produce:

- **Raw Audio:** The original file in its native format, not a converted copy.
- **Comparison Samples:** Any known-authentic speech used for baselining.
- **Model Details:** For machine learning methods, the specific model architecture, version number, and training dataset summary.

- **Configuration:** All parameters used during the analysis (e.g., window size, thresholds, seeds).
- **Underlying Data:** Intermediate outputs, such as anomaly maps or raw score vectors, that support the conclusion.⁴⁹

This level of disclosure ensures that the opposing expert can independently validate the findings, transforming the battle of experts from a distorted feedback loop into a scientific inquiry.

4. *CROSS-EXAMINING THE “BLACK BOX”: A PRACTITIONER’S GUIDE*

When facing an AI expert, counsel and courts should probe the “black box” nature of the tools used. The 2016 PCAST Report emphasized that forensic validity requires empirical testing of the *method* as it is applied.⁵⁰ For AI, this means asking:

- “Did you test on compressed data?” Many commercial detectors claim 99% accuracy on CD-quality audio but fail on Zoom or body-cam audio. If the expert hasn’t validated on the case-specific codec, their error rate is unknown. Lossy compression (e.g., AAC, MP3) often discards the high-frequency information where neural vocoder artifacts reside, blinding many detectors.⁵¹
- “What is the false positive rate on this dialect?” If the defendant speaks AAVE or has a non-native accent, and the model was trained on LibriSpeech (mostly white, native English speakers), the false positive risk skyrockets.⁵²
- “Is the model open-source or proprietary?” If proprietary, has it been independently audited? “Trust me, it’s a trade secret” is incompatible with Sargon’s reliability requirement.
- “Did you use a threshold?” A raw score of “0.8” means nothing without a defined operating point. Did the expert choose a threshold that minimizes false alarms, or one that catches every possible fake?

These questions strip away the AI magic and reveal the underlying statistical reality, allowing the trier of fact to assign appropriate weight.

5. *THE EXPERT’S DUTY TO THE COURT*

In the age of AI, the expert’s duty to the court, to provide impartial, reliable assistance, is paramount. This means resisting the pressure to offer binary “fake/real” certainties when the science supports only probabilities.

- **Transparency:** Experts must explicitly state what their methods *cannot* do. If a detector has a 15% false positive rate on telephony audio, the court must know.
- **Conservatism:** When methods disagree (e.g., artifacts say fake, prosody says real), the expert should report the ambiguity, not suppress the conflicting signal.

- **Education:** The expert's role is to teach the trier of fact how to interpret the data, including the concept of the "liar's dividend" and the limits of detection.

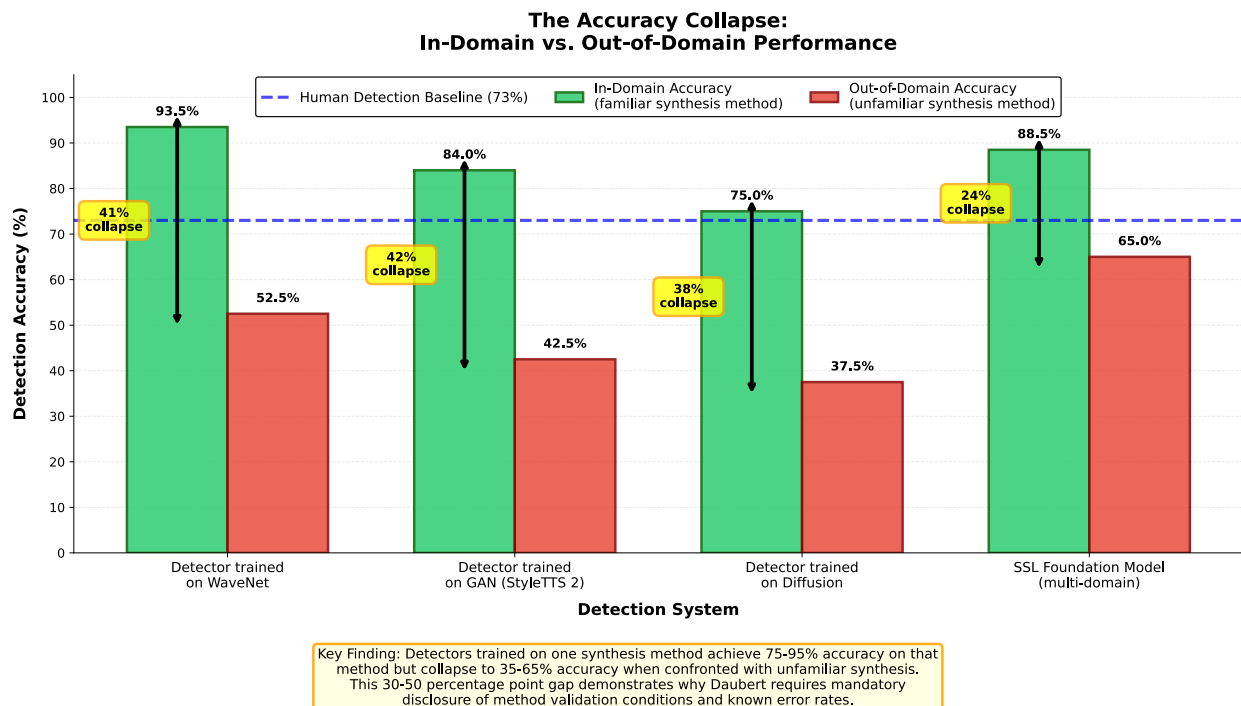
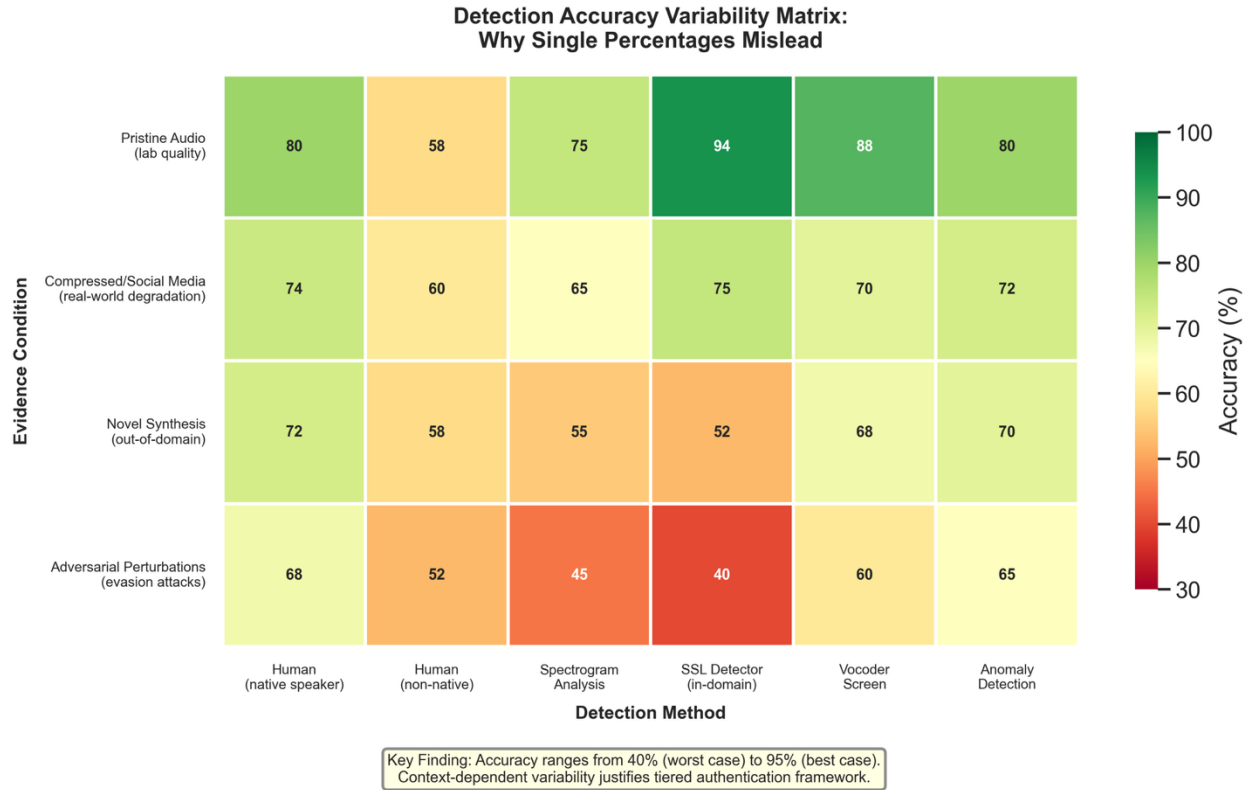
By enforcing these duties, courts ensure that forensic audio analysis remains a tool for truth, not a weapon of confusion.

VI. CONSTITUTIONAL COMPATIBILITY

- **Confrontation Clause:** Authentication addresses reliability, not the truth of statements.⁵³ The defendant retains full cross-examination on methods.
- **Due Process:** Symmetric application benefits indigent defendants by requiring the state to authenticate high-risk evidence and enabling presumptive expert funding.⁵⁴
- **First Amendment:** The Framework governs evidentiary use, not the creation of synthetic speech; it targets fraud, not expression.

VII. CONCLUSION

When synchronized sound first entered the cinema, it marked a permanent break from the past. The silent era was over; a new, more complex reality had begun. We are now living through a similar revolution. The arrival of generative audio is not an incremental change but a fundamental one, forcing us to question the evidence of our own direct perceptions. The challenge this poses to the justice system is profound, reverberating the disruptive promise that we "ain't heard nothing yet."⁵⁵ To navigate this new reality, we cannot rely on the tools of antiquity, and absent this tiered authentication, the economic incentives will favor bad-faith deepfake allegations. The Graduated Framework offers a path forward – a proportional, scientifically-grounded approach designed for the age of AI. It provides the modern tools necessary to distinguish authentic speech from its synthetic shadow, ensuring that as technology speaks in ever more multifarious methods, the law retains its ability to listen, to understand, and to distinguish the truth.



¹Sarah Barrington et al., *Single and Multi-Speaker Cloned Voice Detection: From Perceptual to Learned Features*, arXiv:2307.07683 (July 2023).

²Cal. Evid. Code §§ 1401, 1416.

-
- ³Kevin Warren et al., “Better Be Computer or I’m Dumb”: A Large-Scale Evaluation of Humans as Audio Deepfake Detectors, PROC. 2024 ACM CCS, 2696–2710 (2024).
- ⁴Massimiliano Todisco et al., *ASVspoof 2019: Future Horizons in Spoofed and Fake Audio Detection*, PROC. INTERSPEECH 2019, 1008–12 (2019); Alexei Baevski et al., *wav2vec 2.0: A Framework for Self-Supervised Learning of Speech Representations*, NEURIPS 2020.
- ⁵*People v. Kelly*, 17 Cal. 3d 24 (1976); *Sargon Enterprises, Inc. v. Univ. of S. Cal.*, 55 Cal. 4th 747 (2012).
- ⁶*Mendones v. Cushman & Wakefield, Inc.*, No. 23CV028772, Order re Terminating Sanctions (Cal. Super. Ct., Alameda Cnty. Sept. 9, 2025).
- ⁷*Huang v. Tesla, Inc.*, No. 19CV346663, Tentative Ruling (Cal. Super. Ct., Santa Clara Cnty. Apr. 26, 2023); Bobby Allyn, *People Are Trying to Claim Real Videos Are Deepfakes*, NPR (May 8, 2023).
- ⁸Huang, Tentative Ruling (Apr. 26, 2023).
- ⁹Bobby Chesney & Danielle Citron, *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security*, 107 CAL. L. REV. 1753, 1757–58 (2019).
- ¹⁰Mendones, Order re Terminating Sanctions (Sept. 9, 2025).
- ¹¹Huang, No. 19CV346663; Peter Blumberg & Robert Burnson, *Elon Musk Likely Must Give Deposition in Fatal Tesla Autopilot Crash Suit*, BLOOMBERG (Apr. 27, 2023).
- ¹²Huang, Tentative Ruling (Apr. 26, 2023).
- ¹³Allyn, *supra* note 7.
- ¹⁴Huang, Tentative Ruling (Apr. 26, 2023).
- ¹⁵*Id.*
- ¹⁶Warren et al., *supra* note 3.
- ¹⁷*Id.*
- ¹⁸Innocence Project, *DNA Exonerations in the United States* (2023) (reporting eyewitness misidentification in 69% of cases).
- ¹⁹Nicolas M. Müller et al., *Harder or Different? Understanding Generalization of Audio Deepfake Detection*, arXiv:2406.03512 (June 2024).
- ²⁰Junichi Yamagishi et al., *ASVspoof 2021: Towards Spoofed and Deepfake Speech Detection in the Wild*, IEEE/ACM TRANS. AUDIO, SPEECH, LANG. PROCESS. 31, 3079–97 (2023).
- ²¹Chengzhe Sun et al., *AI-Synthesized Voice Detection Using Neural Vocoder Artifacts*, CVPRW 2023, 1128–37; Zaynab Almutairi & Hebah Elgibreen, *A Review of Modern Audio Deepfake Detection Methods*, ALGORITHMS 15(5), 155 (2022).
- ²²Todisco et al., *supra* note 4.
- ²³Müller et al., *supra* note 19.
- ²⁴NATIONAL RESEARCH COUNCIL, STRENGTHENING FORENSIC SCIENCE IN THE UNITED STATES: A PATH FORWARD 7–8 (2009).
- ²⁵Cal. Evid. Code § 352.
- ²⁶Müller et al., *supra* note 19.
- ²⁷Ian Goodfellow et al., *Generative Adversarial Nets*, NEURIPS 2014 (establishing the generator-discriminator minimax game).
- ²⁸Müller et al., *supra* note 19.
- ²⁹*Id.*; Jinyang Wu et al., *SEA-Spoof: Bridging The Gap in Multilingual Audio Deepfake Detection*, arXiv:2509.19865 (Sept. 2025).
- ³⁰Müller et al., *supra* note 19.
- ³¹*Id.*; Jinyang Wu et al., *SEA-Spoof: Bridging The Gap in Multilingual Audio Deepfake Detection*, arXiv:2509.19865 (Sept. 2025).
- ³²Nicholas Carlini & David Wagner, *Audio Adversarial Examples*, arXiv:1801.01944 (2018).
- ³³Sun et al., *supra* note 21.
- ³⁴Baevski et al., *supra* note 4; Sanyuan Chen et al., *WavLM: Large-Scale Self-Supervised Pre-Training for Full Stack Speech Processing*, IEEE J. SEL. TOP. SIGNAL PROCESS. 16, 1505–18 (2022).

³⁵Yuankun Xie et al., *An Efficient Temporary Deepfake Location Approach Based Embeddings for Partially Spoofed Audio Detection*, arXiv:2309.03036 (Sept. 2023).

³⁶*Sargon*, 55 Cal. 4th at 769–71.

³⁷Cal. Evid. Code §§ 352, 402, 1401; Cal. Civ. Proc. Code § 2019.030.

³⁸Chesney & Citron, *supra* note 9.

³⁹Huang, Tentative Ruling (Apr. 26, 2023).

⁴⁰MCCORMICK ON EVIDENCE § 231 (8th ed. 2020) (noting the rule's purpose to prevent fraud and error).

⁴¹Cal. Evid. Code §§ 1401, 1416.

⁴²NATIONAL RESEARCH COUNCIL, *supra* note 24, at 21–23.

⁴³Cal. Evid. Code § 730.

⁴⁴*United States v. Vayner*, 769 F.3d 125, 131 (2d Cir. 2014).

⁴⁵*Sargon*, 55 Cal. 4th at 769–72.

⁴⁶Cal. R. Ct. 3.722, 3.727; Cal. Evid. Code § 402.

⁴⁷Cal. R. Ct. 3.722, 3.727; Cal. Evid. Code § 402.

⁴⁸*People v. Kelly*, 17 Cal. 3d 24 (1976); *Sargon*, 55 Cal. 4th 747.

⁴⁹Cal. Civ. Proc. Code §§ 2031.010 et seq.; Cal. Penal Code § 1054.1.

⁵⁰PRESIDENT'S COUNCIL OF ADVISORS ON SCI. & TECH., FORENSIC SCIENCE IN CRIMINAL COURTS: ENSURING SCIENTIFIC VALIDITY OF FEATURE-COMPARISON METHODS (2016) (emphasizing empirical testing of validity).

⁵¹Müller et al., *supra* note 19 (noting demographic bias in training data); *see also* Borrelli et al., *Synthetic Speech Detection: A Survey*, arXiv:2302.03823 (2023) (discussing compression robustness).

⁵²Müller et al., *supra* note 19 (noting demographic bias in training data).

⁵³Cal. Evid. Code § 1401.

⁵⁴U.S. Const. amends. V, VI.

⁵⁵*The Jazz Singer* (Warner Bros. 1927) marked the first feature-length motion picture with synchronized dialogue. Al Jolson's famous line, "You ain't heard nothin' yet," signaled the end of the silent era and the beginning of the "talkies," a technological disruption parallel to the generative AI revolution.