

Training Data Transparency Disclosure

Developer: Eleven Labs Inc. (“**ElevenLabs**”)

Last Updated: January 1, 2026

Effective Date: January 1, 2026

Jurisdictional Applicability: California (AB 2013) and other jurisdictions where training data transparency disclosures are required or otherwise provided.

This disclosure (“**Disclosure**”) is provided to promote transparency regarding the data used to train ElevenLabs’ generative artificial intelligence (“**AI**”) systems and services. It is intended to satisfy applicable training data transparency and disclosure requirements, including California Civil Code §§ 3110-3111 (AB 2013), and to provide similar transparency information to users in other jurisdictions where comparable transparency obligations apply.

Model Families Covered by This Disclosure

ElevenLabs organizes its generative AI systems into model families based on modality, intended function, and training data characteristics. Each model family described below constitutes a distinct “generative AI system or service” for purposes of AB 2013 and analogous transparency frameworks.

1. Text-to-Speech (TTS) Models

Operational domain: AI audio – synthetic speech generation.

Primary function: Generate spoken audio from text input.

Model Versions:

ElevenLabs’ TTS systems may be deployed in multiple model versions or iterations (e.g., v1, v2, v3). These versions reflect ongoing development and may differ in architecture, scale, and training data composition over time.

Training Data Sources

Training data for TTS models may include, across different model versions, a combination of:

- Licensed voice recordings, including across multiple languages, accents, and speaking styles;
- Publicly available data, including lawfully accessible speech and text data;
- Proprietary data, including datasets created, commissioned, or curated by ElevenLabs; and

- User-provided data, where users have permitted such use or where ElevenLabs otherwise has a lawful basis to use the data for training or improvement purposes.

Not all training data sources are used in every model version, and the relative contribution of each source may vary by model version.

Dataset Scale and Characteristics

- Datasets may include large numbers of voice samples across multiple languages, accents, and speaking styles.
- Data points may include voice recordings and associated metadata (e.g., language, speaker, or recording conditions).

Intellectual Property and Personal Data

Training data may include copyrighted material and personal information, including voice data. Voice data is used subject to applicable privacy, consent, and data protection requirements.

2. Voice Models (Voice Conversion and Voice Cloning)

Operational domain: AI audio – voice transformation and voice cloning.

Primary function: Transform or generate speech using specific voice characteristics.

Training Data Sources

Training data for Voice models may include a combination of:

- Licensed voice recordings, including across multiple languages, accents, and speaking styles;
- Proprietary data, including datasets created, commissioned, or curated by ElevenLabs; and
- User-provided data, where users have permitted such use or where ElevenLabs otherwise has a lawful basis to use the data for training or improvement purposes.

Dataset Scale and Characteristics

- Datasets may include voice samples across multiple languages, accents, and speaking styles.
- Data points may include voice recordings, and associated metadata (e.g., language, speaker, or recording conditions).

Intellectual Property and Personal Data

Training data may include copyrighted material and personal information, including voice data. Voice data is used subject to applicable privacy, consent, and data protection requirements.

3. Speech-to-Text (STT) Models

Operational domain: AI audio – transcription of speech into text.

Primary function: Convert spoken audio into written text or structured textual representations.

Training Data Sources

Training data for STT models may include a combination of:

- Licensed speech recordings paired with corresponding text transcriptions;
- Publicly available data, including lawfully accessible speech and transcription datasets;
- Proprietary data, including datasets created, commissioned, or curated by ElevenLabs; and
- User-provided data, where users have permitted such use or where ElevenLabs otherwise has a lawful basis to use the data for training or improvement purposes.

Dataset Scale and Characteristics

- Training datasets may include large volumes of audio data and corresponding text, expressed as ranges or estimates.
- Data points may include audio recordings, text transcriptions, and associated metadata (e.g., language, speaker, or recording conditions).

Intellectual Property and Personal Data

Training datasets may include copyrighted content and personal information, including voice data. Data is used in accordance with applicable privacy and data protection laws.

4. Music Models

Operational domain: AI audio – music audio generation.

Primary function: Generate or transform music.

Training Data Sources

Training data for Music models may include a combination of:

- Licensed music recordings;
- Proprietary datasets, including internally created, commissioned audio, or internally labeled datasets derived from or based on licensed music recordings; and
- Synthetic audio, generated to supplement training or evaluation.

Dataset Scale and Characteristics

- Dataset sizes may be expressed as ranges or estimates.
- Data points may include music audio recordings and associated metadata.

Intellectual Property

Training data for these models may include content protected by copyright or other intellectual property rights and is used in accordance with applicable licensing terms.

5. Non-Speech Audio Models

Operational domain: AI audio – non-speech audio generation.

Primary function: Generate or transform other non-speech audio content.

Training Data Sources

Training data for Non-Speech audio models may include a combination of:

- Licensed non-speech audio recordings, including sound effects;
- Proprietary datasets, including internally created or commissioned audio; and
- Synthetic audio, generated to supplement training or evaluation.

Dataset Scale and Characteristics

- Dataset sizes may be expressed as ranges or estimates.
- Data points may include audio recordings (e.g., sound effects or ambient audio) and associated metadata.

Intellectual Property

Training data for these models may include content protected by copyright or other intellectual property rights and is used in accordance with applicable licensing terms.

Substantial Modifications

ElevenLabs updates its training datasets and models on an ongoing basis. This Disclosure will be updated upon any substantial modification, including material changes to training data sources, dataset composition, or model performance characteristics, prior to or concurrent with public availability in California and other applicable jurisdictions. Changes in training data composition across successive versions within an existing model family do not, by themselves, constitute a new model family unless they result in a material change to the system's modality or primary function.

Limitations

This Disclosure is intended to provide meaningful transparency while protecting confidential business information, trade secrets, and security-sensitive details. Certain information is therefore provided in summarized or estimated form, as permitted by law. References to user-provided data in this Disclosure are subject to ElevenLabs' applicable privacy notices and data protection obligations. This Disclosure is not intended to create contractual commitments or expand ElevenLabs' obligations beyond those required by applicable law.

Contact Information

Questions regarding this disclosure may be directed to:

Eleven Labs Inc.

Legal & Compliance Team

Email: legal@elevenlabs.io