

Lyria 3.5

Model Card

Model Cards are intended to provide developers with essential, summarized information on models, including overviews of known limitations and mitigation approaches. Model cards may be updated from time to time; for example, to include updated evaluations as the model is improved or revised. The model card below covers Lyria 3.5 and subsequent versions.

Published: July 2026

Model Information

Description: Lyria 3.5 is a music generation system capable of synthesizing high-quality audio from a text prompt.

Inputs: Text

Outputs: Audio (music), text (lyrics)

Architecture: Lyria 3.5 utilizes [latent diffusion](#), applied to temporal audio latents.

Model Data

Training Dataset: Lyria 3.5 was trained on audio data. Audio datasets were annotated with text captions at different levels of detail.

Training Data Processing: Data filtering and preprocessing included techniques such as deduplication, safety filtering in-line with [Google's commitment to advancing AI safely and responsibly](#), and quality filtering to mitigate risks and improve training data reliability and compliance. Once data is collected, it is cleaned and preprocessed to make it suitable for training.

Implementation and Sustainability

Hardware: Lyria 3.5 was trained using [Google's Tensor Processing Units \(TPUs\)](#). TPUs are specifically designed to handle large-scale computation involved in training and serving generative AI models, increasing training throughput and efficiency considerably when compared to CPUs. TPUs come with large amounts of high-bandwidth memory, allowing for the handling of large models and batch sizes during training, which can lead to better model quality. TPU Pods (large clusters of TPUs) also provide a scalable solution for handling the growing complexity of large foundation models. Training can be distributed across multiple TPU devices for faster and more

efficient processing.

The efficiencies gained through the use of TPUs are aligned with Google's [commitment to operate sustainably](#).

Software: Training was done using [JAX](#) and [ML Pathways](#).

Distribution

Lyria 3.5 is distributed in the following channels; respective documentation shared in line:

- [Google Flow Music](#)

Our models are available to downstream providers via an application program interface (API) and subject to relevant terms of use. There is no required hardware or software to use the model.

Evaluation

Approach: Lyria 3.5 was evaluated across a range of benchmarks, including human and automated evaluations. In close collaboration with music experts, we have developed an evaluation framework that tests the model's ability to represent concepts (such as genres, moods, instruments) in- and out-of-distribution across a broad range of curated user prompts. We sampled several music generation models for comparison, including Lyria 2.

We focus on a range of quality dimensions including music quality and aesthetics, vocal quality, audio fidelity, and prompt adherence.

Results: Lyria 3.5 improved significantly compared to Lyria 2 on audio fidelity, highlighting the technical quality and clarity of the generated audio. With lyrics, Lyria 3.5 demonstrates better prompt adherence, following both simple and more complex instructions more accurately.

Intended Usage and Limitations

Benefit and Intended Usage: Lyria 3.5 is Google's most capable music generation model to date. Lyria can be used to generate high-quality music in a wide range of genres and styles.

Acceptable Usage: [Google's Generative AI Prohibited Use Policy](#) applies to uses of the model in accordance with the applicable terms of service. Additionally, the model should not be integrated

into certain systems (also found in [Google's Generative AI Prohibited Use Policy](#)), including those that: (1) engage in dangerous or illicit activities, or otherwise violate applicable laws or regulations, (2) violate the rights of others, including privacy and intellectual property rights, (3) compromise the security of others' or Google's services, (4) engage in sexually explicit, violent, hateful, or harmful activities, (5) engage in misinformation, misrepresentation, or misleading activities.

Ethics and Safety

Evaluation Approach: Lyria 3.5 was developed in partnership with internal safety and responsibility teams. A range of evaluations and red-teaming activities were conducted to help improve the model and inform decision-making. These evaluations and activities align with [Google's AI Principles](#) and [responsible AI approach](#). Evaluation types included but were not limited to:

- **Training/Development Evaluations** including automated and human evaluations carried out continuously throughout and after the model's training, to monitor progress and performance
- **Human Red Teaming** conducted by specialist teams who sit outside of the model development team deliberately trying to spot weaknesses and ensure the model adheres to safety policies and desired outcomes
- **Ethics & Safety Reviews** were conducted ahead of the model's release

In addition to the evaluations above, system-level safety evaluations and reviews are run within the context of specific applications that models are deployed within.

Risks & Mitigations: Safety and responsibility was built into Lyria 3.5 throughout the training and deployment lifecycle, including pre-training, post-training, and product-level mitigations. Mitigations include, but are not limited to:

- dataset filtering;
 - conditional pre-training;
 - supervised fine-tuning;
 - reinforcement learning from human and critic feedback;
 - safety policies and desiderata;
 - product-level mitigations such as safety filtering and applying [SynthID](#) watermarking.
-