

Audion Voice AI Studio

[Русский](#) · [User Guide](#)

Contents

- [Why It Exists](#)
- [What Has to Be Installed Separately](#)
- [Editions](#)
- [The Principle](#)
- [What It Can Do](#)
- [Next](#)
- [Technical Reference](#)
 - [FFmpeg and the NVIDIA Driver](#)
 - [Portability](#)

Transcribing recordings, live dictation, cleaning up text, and exporting working notes. The extended edition — for a workstation with an NVIDIA card.

Why It Exists

Transcription is needed in two entirely different situations.

Quickly, right now. Dictate a paragraph, transcribe a ten-minute recording, get the text immediately. Response time is what matters, and a cloud service fits best.

In bulk, offline. An hour-long meeting, a whole day of dictaphone recording, a folder of files. Here autonomy matters more: the recording should not leave the machine, and the cost should not grow with every hour.

This edition leans on the second. On a machine with a graphics card, local engines stop being a compromise: an hour of recording is processed in minutes, and the recording never leaves the computer.

What Has to Be Installed Separately

Model weights are not part of the distribution — about 7 GB: the GigaAM cache and whisper.cpp with the Turbo and Large V2 models; with the speaker separation layer (NVIDIA only) about 10.6 GB.

The weights carry their own licences, separate from the program's, and the terms for speaker separation are accepted **personally, under your own account** — so the user downloads them.

Until the models are installed, transcription will not start: the window opens, but there is nothing to recognise with.

On the first start the app offers to download what is missing. The "Download models and engines" window lists the modules with their download size: GigaAM, whisper.cpp CUDA with the Turbo model, the Large V2 model and, on an NVIDIA machine, the speaker-separation layer; about 10.6 GB together. "Download and install" installs them one after another; progress is shown on the Maintenance page, and once everything is in place every mode there is green. "Later" postpones the question until the next start, "Don't ask again" hides the window for good. Modules can still be installed by hand on the same page.

Editions

edition	for what	engines
Live	laptop, everyday work	cloud services, GigaAM, whisper.cpp
Studio	workstation with an NVIDIA card	the same plus CUDA and speaker separation

Live remains the main version for ordinary work. Studio adds accelerated local engines, large models, and speaker separation.

The Principle

Reliable first, elegant second. Batch transcription behaves predictably, and everything else rests on it. An overlay on top of other windows and instant on-the-fly parsing were both tried — and both broke the window itself. Heavy work is not hung on the thread that draws the interface.

What It Can Do

Transcribing files and folders, live dictation, separating a recording by speaker, cleaning up transcribed text, exporting into working notes, a tray icon for quick access, resetting the application to its initial state.

Next

- [User Guide](#) — step by step, engines, formats.
-

Technical Reference

FFmpeg and the NVIDIA Driver

Every FFmpeg build is compiled against a particular version of the hardware-encoding headers, and each demands its own minimum driver. The newest build on an old driver does not accelerate anything — it breaks the hardware path. So the build is chosen to match the driver, not by version number.

Portability

Application state lives in the project folder, not in the system. A separate command returns the application to its initial state without touching the downloaded models.