

Getting started with voices

Studios without floor time, mock-ups, previews

CONTENTS

1. Cast the voices and generate an episode · 20 minutes
2. Adjust the lips on screen · 5 minutes of setup, then the processing

Cast the voices and generate an episode

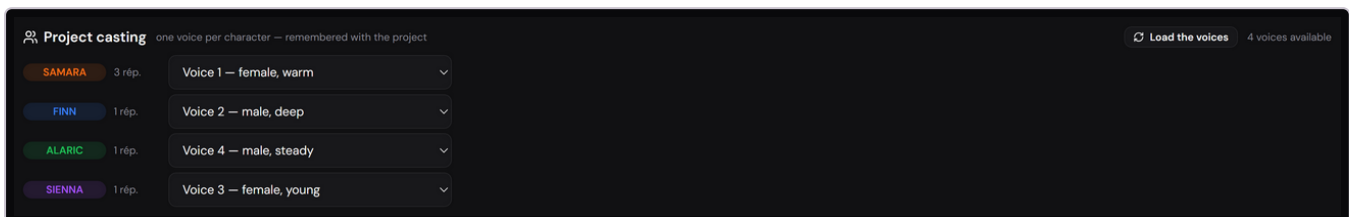
The dubbing director who wants to hear a mock-up before calling the actors in.

20 minutes

01 Cast the project

Load the voices pulls the list available on your account, then you assign one voice per character. The casting belongs to the project: it carries across episodes.

You connect **your own credentials**. Costs, quotas and voice rights stay yours, which is also what leaves you in control of the provider.

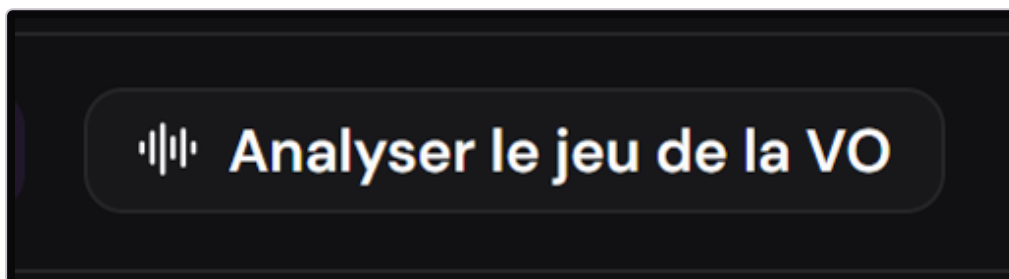


The project casting: one voice per character, with each one's line count.

02 Listen to the original performance

Analyse the original performance measures, line by line, the intensity, pitch, rate and intonation of the original actor. Those measurements are then attached to the tag request: the AI stops guessing the emotion from the meaning of the words alone.

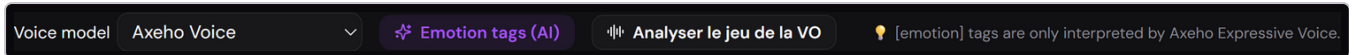
The analysis is **local**, with no network and no cost. Each line is compared to the average **of the same character**: a bass voice is not "low", it is normal for that actor.



The button that analyses the original performance.

03 Set the performance intent

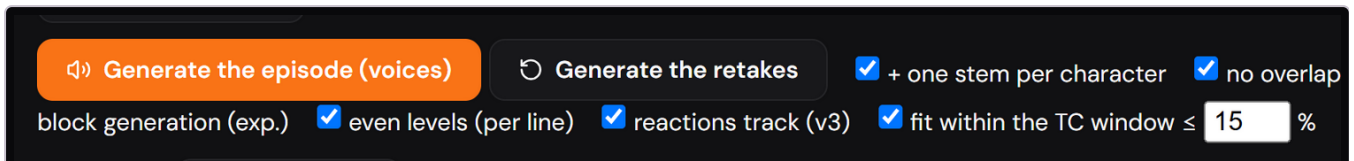
Emotion tags suggests performance markers line by line. You can also insert them by hand wherever the reading has to change.



The emotion tags button and the analysis of the original performance.

04 Generate the episode

Generation produces clips **aligned to the project's timecodes**, plus separated stems. Each run is dated and kept: you can compare two versions instead of overwriting the last one.

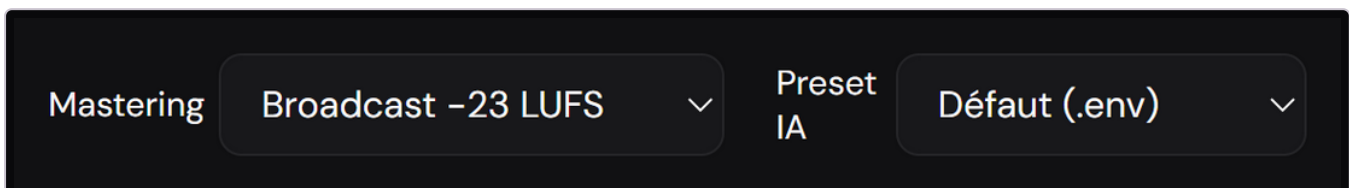


The episode generation button and the track options.

05 Master for the broadcaster

Four loudness targets: **broadcast -23**, **streaming -27**, **web -16**, **micro-drama -15 LUFS**. It applies to the full mix and to every stem, and the broadcaster no longer has to reject the delivery over a level that is not theirs.

The mastered file is written **next to** the raw one, never over it: you compare the two, and keep the one you send.



The mastering target and processing preset menus.

Adjust the lips on screen

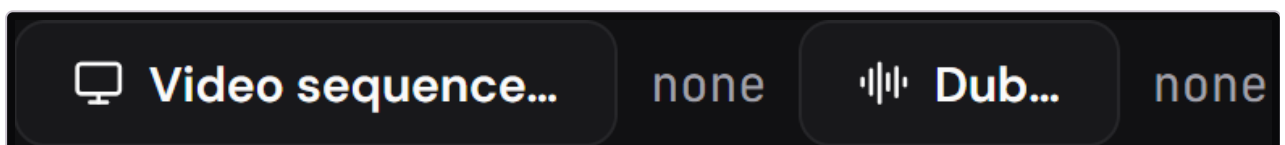
An advert, a corporate film, a training module: the picture has to match the dub.

5 minutes of setup, then the processing

01 Point at the sequence and the dub

A **video sequence** and the **dub** that goes with it. The feature works on a shot, not on a whole episode.

Not for fiction. The actor's performance is the product; distorting it damages the very thing you are paid to deliver. Axeho's reference sync remains the one in the text, words chosen to land where the mouth closes.



Picking the video sequence and the dub file.

02 Process only the part that matters

Bound the passage in seconds. Less data sent, so cheaper, and a better result: the service works on what matters instead of crossing the whole shot.

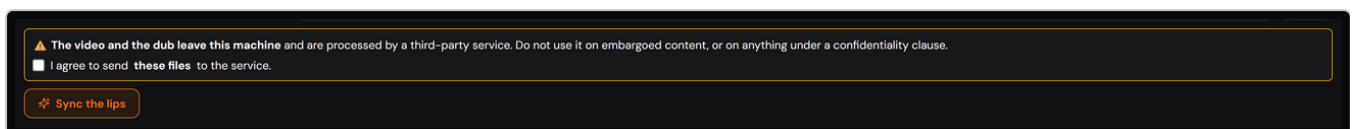


The start and end bounds of the excerpt to process, in seconds.

03 Accept the upload, then start

The video and the dub leave this machine and are processed by a third-party service. The box is ticked each time, for **those files**: it is a decision you take, not a setting you forget.

Everything else on this page works on the machine. This option is the only one that leaves, which is why it asks for explicit consent.



The third-party upload warning, its consent box and the start button.

This tutorial online : <https://axeho.studio/en/tutorials-voices.html>