

Macローカル議事録AI 導入チェックリスト&トラブルシューティング早見表

公式LINE読者特典 | 音声を1秒も外に出さない議事録自動化 (whisperX + Ollama) | aki studio

1. 導入前チェックリスト（始める前に全部☑）

- Apple Silicon搭載Mac（M1以降）である（りんごマーク → 「このMacについて」で確認）
- メモリ16GB以上ある
- ディスク空き容量が12GB以上ある
- Homebrewがインストール済み（`brew --version` が通る）
- HuggingFaceのアカウントを作成した（無料・話者分離に必須）
- pyannote/speaker-diarization-community-1 のページで利用規約に同意した
- HuggingFaceのアクセストークンを発行し `~/.config/whisperx/hf_token` に保存した
- Ollamaは cask 版（`brew install --cask ollama-app`）で入れた（formula版はNG）
- 最後に `bash ~/gijiroku_local/healthcheck.sh` を実行し、5項目すべて🟢になった

🔧 困ったら最初にこれ: `bash ~/gijiroku_local/healthcheck.sh` を実行すると、ffmpeg/whisperX/HFトークン/Ollama/ディスク空きの5項目を一括点検し、壊れている箇所と対処法を表示します。macOSアップデートや `brew upgrade` の後に動かなくなったときも、まずこれです。

2. インストールで詰まったら（症状 → 原因 → 対処）

症状	原因	対処
<code>ollama: command not found</code>	Ollama.app未起動 or PATH未反映	Ollama.appを一度起動 → ターミナルを開き直す
whisperXのpip installが失敗する	Pythonが3.13以降	必ず <code>python3.12 -m venv</code> で仮想環境を作る（3.13以降は依存ライブラリ非対応）
torch関連のエラーで止まる	pipが古い	<code>pip install --upgrade pip</code> してから再実行
ffmpegエラー（dylib not found等）	brewの他パッケージ更新でリンク切れ	<code>brew reinstall ffmpeg</code> （タップ版でなくcore版を使う）

3. 実行時のエラー早見表

症状	原因	対処
「Ollamaが起動していません」	サーバ停止中	メニューバーにラマのアイコンがあるか確認。なければOllama.appを起動
401 Unauthorized (HuggingFace)	トークン未設定・タイプミス	<code>cat ~/.config/whisperx/hf_token</code> で中身確認。前後の空白・改行に注意

403 Forbidden / gated repo	モデルの利用規約に未同意	HuggingFaceで speaker-diarization-community-1 と segmentation-3.0 の両方に同意する
処理が異常に遅い・Macが固まる	メモリ不足でスワップ発生	他の重いアプリ（ブラウザのタブ大量等）を閉じる。WHISPERX_MODEL=small にする
要約が途中で切れる・浅い	会議が長くコンテキスト超過	スクリプトの分割要約（9,000字単位）が動いているか確認。それでも浅ければ会議を前後半に分けて実行
要約に中国語の漢字が混ざる	qwen系モデルの癖	スクリプト内の変換テーブル（KANJI_FIX）対処済み。それ以外の字が出たらテーブルに追記
「ご視聴ありがとうございました」など言っていない文が入る	Whisperの幻覚（無音・BGM区間で定型句を捏造する仕様）	スクリプトの幻覚フィルタが自動除去（除去内容は removed_hallucinations.diff で確認可）。別パターンの捏造文を見つけたらフィルタの正規表現に追記
登録した人名が文中に大量に湧く	turbo系モデル+辞書の副作用（実測で215回捏造）	モデルは large-v3（既定）のまま使う。turboは辞書と併用しない

4. 精度を上げる6つのコツ

コツ	内容
① 話者数を指定する	参加人数がわかっているなら gijiroku.sh 音声.m4a 2 4 のように最小・最大を渡す（実人数±1が目安）。話者分離の精度が大きく上がる
② 固有名詞辞書を置く	実行フォルダに glossary.txt（1行1語で人名・社名・専門用語）を置くだけで認識がその表記に寄る。実測で登録5氏名すべて正表記化。毎回同じ会議なら使い回す
③ 話者を実名にする	一度実行して誰がSPEAKER_00かを確認したら、SPEAKERS="SPEAKER_00=田中,SPEAKER_01=佐藤" を付けて再実行。議事録も全文も実名表記になる。定例会議なら対応表は毎回同じになりやすい
④ マイクは部屋の中央に	話者分離は声の音響的特徴で判定する。全員から等距離だと分離しやすい。スマホ録音なら机の中央に置く
⑤ 声かぶりを減らす	同時発話の区間は分離も文字起こしも崩れる。会議の進め方の問題だが効果は一番大きい
⑥ 議事録は「初稿」と考える	ローカルAIの要約は完璧ではない。決定事項とToDoだけ人間が3分見直す運用が最も費用対効果が高い

 **モデル選びの実測データ（42分の実会議・M4/16GB）** : large-v3=41分・精度◎ / medium=34分・内容語の誤り増 / turbo=35分・辞書と併用不可。処理時間の大半は話者分離側なので、**モデルを小さくしても大して速くなりません。既定のlarge-v3のままだが正解**です（メモリ不足で固まる機種のみ WHISPERX_MODEL=small）。

💡 **豆知識:** モデル（約6GB）は議事録生成のときだけメモリに載り、数分使わないと自動で解放されます。Ollamaを常駐させてもMacは重くなりません。

5. それでも解決しないときは

このチェックリストで解決しない場合や、「そもそも構築から任せたい」という方向けに、画面共有によるリモート構築代行（動作確認・2週間サポート込み）を承っています。公式LINEからお気軽にご相談ください。

aki studio | AI×副業の実験記録

公式LINE: <https://lin.ee/vrsG2v38> | Instagram: @ai_kasegu_jikken | note: aki__studio

※本資料の再配布はご遠慮ください（LINE読者特典です）