

08-minuspod.md

kstack: book: Centerpoint Home
Lab chapter: Utilities page:

MinusPod tags: [minuspod,
podcasts, transcription, whisper,
ai, cuda]

Overview

MinusPod is an AI-powered podcast manager and transcriber. It manages podcast feeds, downloads episodes, and transcribes audio using a local Whisper model. Summaries and notes are generated via an LLM (Qwen3:14b) through the Ollama API. The container runs with NVIDIA GPU access (RTX 5080) for accelerated Whisper inference. Internal-only access. Episode data is retained for 180 days.

Access

Type	URL	Auth
Internal	<code>https://minuspod.home.local</code>	MinusPod own auth (Step-CA TLS)

No external route — LAN access only.

Configuration

Image: `ttlequals0/minuspod:2.1.9`
Runtime: CUDA 12.9 / NVIDIA runtime (RTX 5080)

Key Environment Variables

Variable	Value / Notes
<code>BASE_URL</code>	<code>https://minuspod.home.local</code>
<code>WHISPER_BACKEND</code>	<code>local</code>
<code>WHISPER_DEVICE</code>	<code>cuda</code>
<code>WHISPER_MODEL</code>	<code>medium</code>
<code>LLM_PROVIDER</code>	<code>ollama</code>
<code>OPENAI_BASE_URL</code>	<code>https://ollama.home.local/v1</code>
<code>OPENAI_MODEL</code>	<code>qwen3:14b</code>
<code>OPENAI_API_KEY</code>	REDACTED
<code>RETENTION_PERIOD</code>	<code>4320</code> hours (180 days)
<code>MINUSPOD_TRUSTED_PROXY_COUNT</code>	<code>1</code>
<code>NVIDIA_VISIBLE_DEVICES</code>	<code>all</code>
<code>NVIDIA_DRIVER_CAPABILITIES</code>	<code>compute,utility</code>

Traefik Labels

```
traefik.http.routers.minuspod.rule: Host(`minuspod.home.local`)
traefik.http.routers.minuspod.entrypoints: websecure
traefik.http.routers.minuspod.tls.certresolver: step-ca
traefik.http.services.minuspod.loadbalancer.server.port: 8000
```

Volumes / Bind Mounts

Host Path	Container Path	Purpose
<code>/home/jeeves/dockers/minuspod/data</code>	<code>/app/data</code>	Episode database and downloads

“ Note: The data directory is under `/home/jeeves/dockers/` (not `docker/`) — this is the actual bind mount path as confirmed by `docker inspect`.

AI Integration

MinusPod uses two AI components:

Component	Model	Transport	Purpose
Whisper	medium	Local CUDA inference (RTX 5080)	Audio transcription
Qwen3:14b	qwen3:14b	Ollama API at ollama.home.local/v1	Summaries and notes

`OPENAI_BASE_URL` points to the internal Ollama instance (which exposes an OpenAI-compatible API). The `OPENAI_MODEL` is `qwen3:14b` — this model must be present in Ollama. Verify with `docker exec ollama ollama list`.

Notes / Gotchas

- The Whisper `medium` model on CUDA provides a good balance of speed and accuracy. Upgrade to `large-v3` if accuracy is more important than speed (uses more VRAM). The RTX 5080 has 16GB GDDR7 — ample for either.
- `RETENTION_PERIOD=4320` hours = 180 days. Episodes older than this are automatically removed. Adjust if long-term episode storage is needed.
- `MINUSPOD_TRUSTED_PROXY_COUNT=1` tells MinusPod to trust one level of proxy headers (Traefik) for correct IP and protocol detection.
- The bind mount path is `/home/jeeves/dockers/minuspod/data` (note: `dockers`, not `docker`). This is different from the convention used by most other stacks.
- Ensure `qwen3:14b` is loaded in Ollama before enabling LLM summaries:

```
docker exec ollama ollama pull qwen3:14b
```

Last Updated: 2026-06-17