

Local AI

Einführung, Erfahrungen,
persönlicher Stack

Local AI - Warum?

AI

Jensen Huang says he would be 'deeply alarmed' if his \$500,000 engineer did not consume at least \$250,000 of tokens

KI

Angeblich halbe Milliarde Dollar in Tokens verbrannt

Ein Konzern soll rund 500 Millionen Dollar für Claude ausgegeben haben, weil niemand Nutzungslimits für Mitarbeiter eingerichtet hatte.

31. Mai 2026 um 12:11 Uhr / Michael Linden

US-Regierung erzwingt Abschaltung von Anthropics KI Fable 5 und Mythos 5

Eine Exportdirektive der US-Regierung zwingt Anthropic, Fable 5 und Mythos 5 abzuschalten. Das Unternehmen spricht von einem Missverständnis.

Local AI - Wie?

Local AI - Wie?

Voraussetzungen

Hardware:

- GPUs (Nvidia, AMD)**
- (V)RAM**
- Apple Silicon (Shared Memory)**

Local AI - Wie?

Welche Modelle sind mit meiner Hardware kompatibel?



<https://github.com/AlexsJones/llmfit>



<https://huggingface.co/models>

Login -> Settings -> Hardware



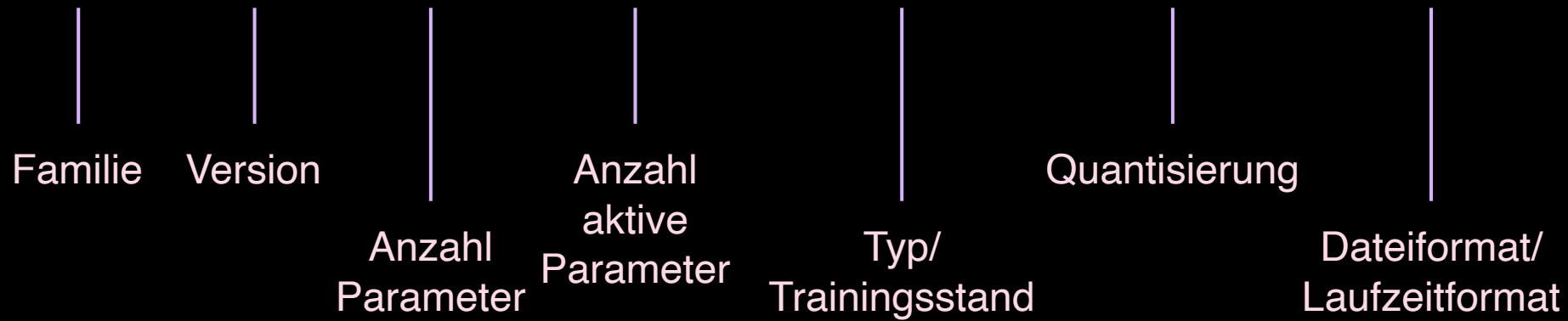
<https://artificialanalysis.ai>

Local AI - Wie?

Welches Modell wofür?

-> Open Weight Bezeichnungen verstehen

qwen3.6:35b-a3b-coding-nvfp4 (MLX)



Local AI - Wie?

Wie kann ich Modelle lokal installieren und ausführen (Inference)?

-> Beispiel: macOS



<https://omlx.ai/>



<https://ollama.com/>

 oMLX
v0.5.7

Dashboard

Models

Settings

Logs

Bench

Chat

 [Logout](#)

DASHBOARD

Serving Stats

Session

All-Time

Clear

Qwen3.6-35B-A3B-MLX-4bit

TOTAL PREFILL TOKENS

7.143

CACHED TOKENS

2.048

CACHE EFFICIENCY

28.7%

AVERAGE SPEED

Prompt Processing (excl. cached)

556.1 tok/s

Token Generation

82.8 tok/s

ACTIVE MODELS

Qwen3.6-35B-A3B-MLX-4bit

idle 59s

~19.36GB obs / 19.95GB est

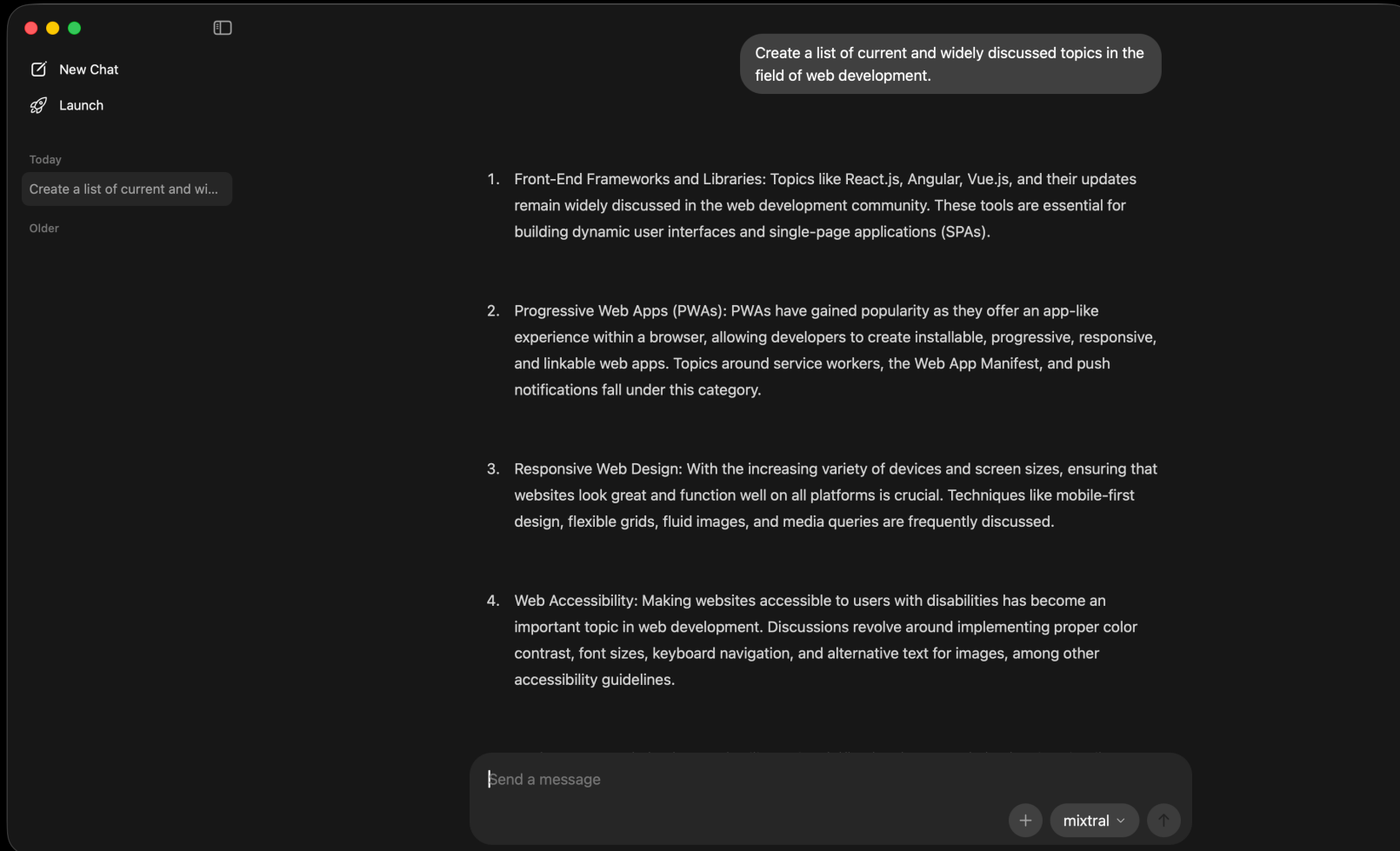
RUNTIME CACHE OBSERVABILITY

SSD

203 MB / 92.0 GB · 2 files

Active base_path

Effective ssd_cache_dir



Local AI - Wofür?

Für welche Aufgaben kann ich Local AI anwenden?

 Claude Code

HERMES
AGENT

 Obsidian

opencode

 Raycast

 Whisper Notes

Danke

Fragen bitte an: Martin Wilmer <martin.wilmer@robole.de>