

O HackTown é um espaço apartidário

⚠ Qualquer manifestação político-partidária identificada antes, durante ou após a apresentação pode resultar na interrupção imediata da atividade e na exclusão do speaker da programação, sem prejuízo das medidas previstas no Termo de Speaker assinado.





HACKTOWN · PALESTRA

Criando seu assistente pessoal com IA local e independente

Ferramentas abertas para montar sua própria IA, peça por peça — com a opção de ligar a nuvem só quando você quiser.

Vitor Casadei

Senior Data Science Manager, CESAR · linkedin.com/in/vcasadei



Por que IA local e independente?



Privacidade

Seus e-mails, sua agenda e suas reuniões não precisam passar pelo servidor de ninguém.



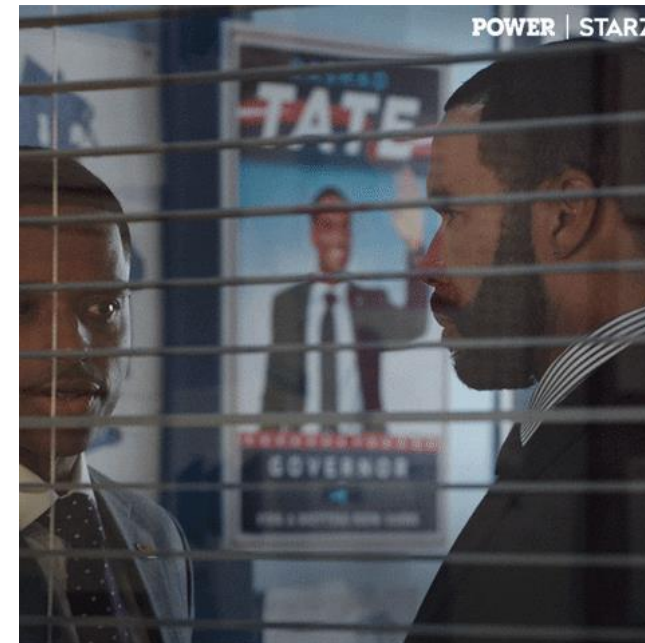
Custo e dependência

Sem internet, sem custo por token e sem fila de espera de um serviço em nuvem.



Aprisionamento a um fornecedor

Você escolhe quando (e se) liga uma nuvem — nunca é obrigado.





Eu já uso uma peça disso todo dia

- Comecei tentando organizar reuniões no OneNote/Obsidian — funcionava, mas eu exportava tudo à mão só pra poder fazer perguntas a uma IA depois.
- Hoje eu gravo com o Meetily: transcrição e resumo 100% no meu PC — o áudio nunca sai da máquina.
- Curiosidade: até a narração desta história que virou os slides eu gravei e transcrevi com o próprio Meetily.
- Funciona tão bem que virou o ponto de partida de uma jornada — Meetily sozinho → AnythingLLM → ODS → OpenClaw — que é exatamente o que vou te mostrar hoje.



Meetily: 100% open source

meetily.ai — github.com/Zackriya-Solutions/meetily · licença MIT, grátis

- Grava, transcreve (Whisper ou Parakeet) e resume reuniões — tudo local, o áudio nunca sai da máquina.
- Windows e Mac; o resumo final roda por um modelo do Ollama ou por uma API à sua escolha.



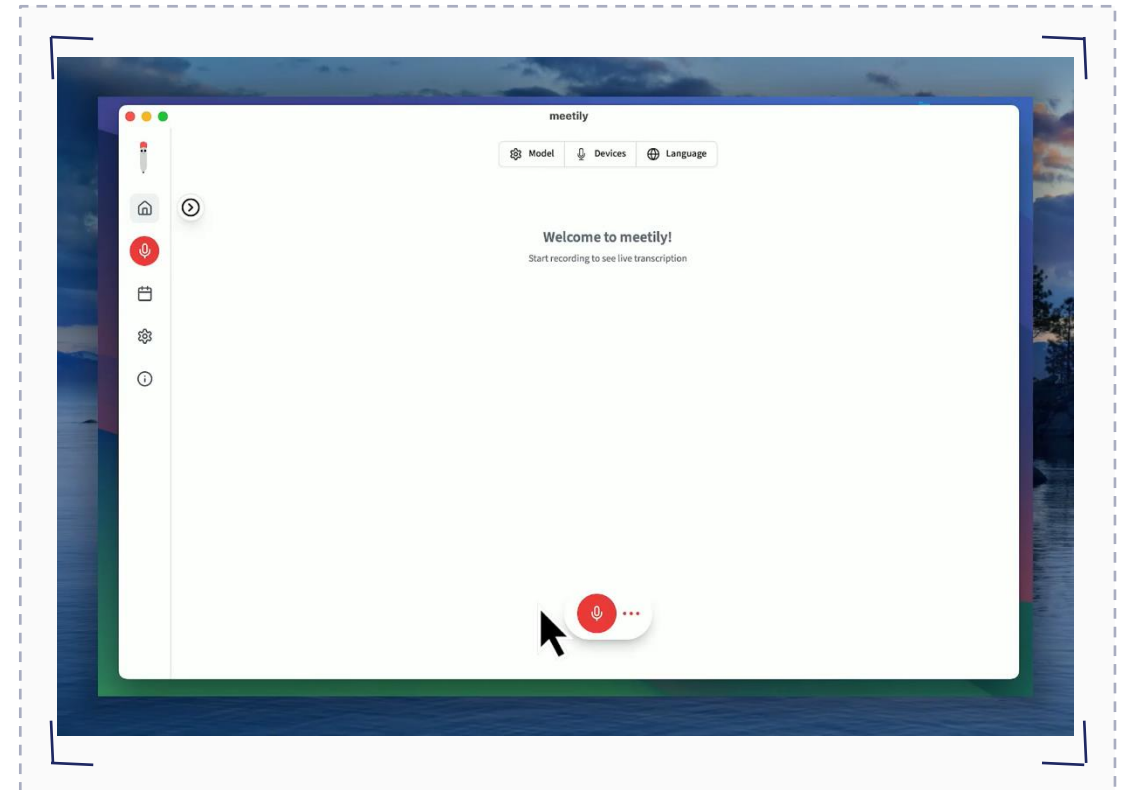
O AnythingLLM já faz isso sozinho

O "Meeting Assistant" dele grava, transcreve e resume nativamente (Mac/Windows) — no caminho AnythingLLM, dá pra nem precisar do Meetily.

COMO INSTALAR

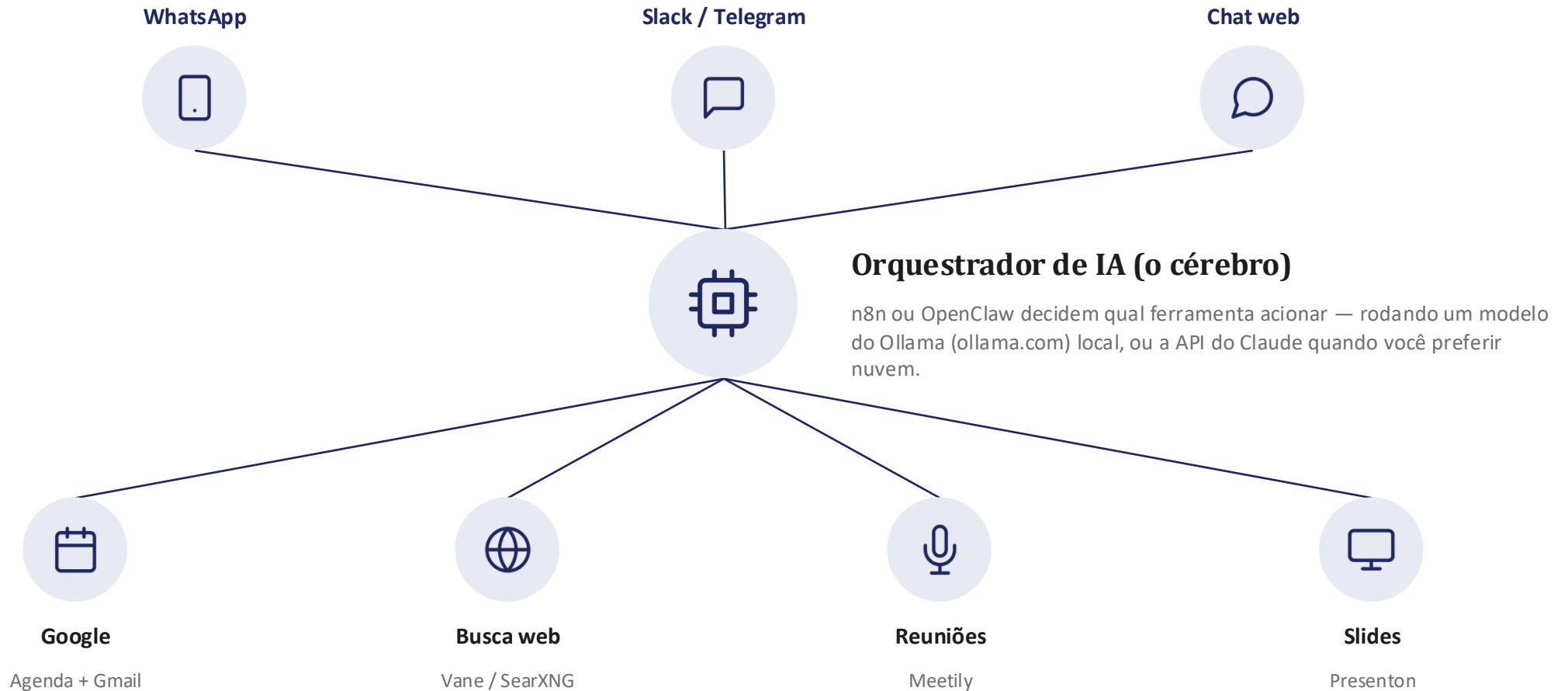
Windows: instalador nas releases do GitHub

Mac: `brew install --cask meetily`





Como as peças conversam entre si





Ollama roda o modelo na sua máquina

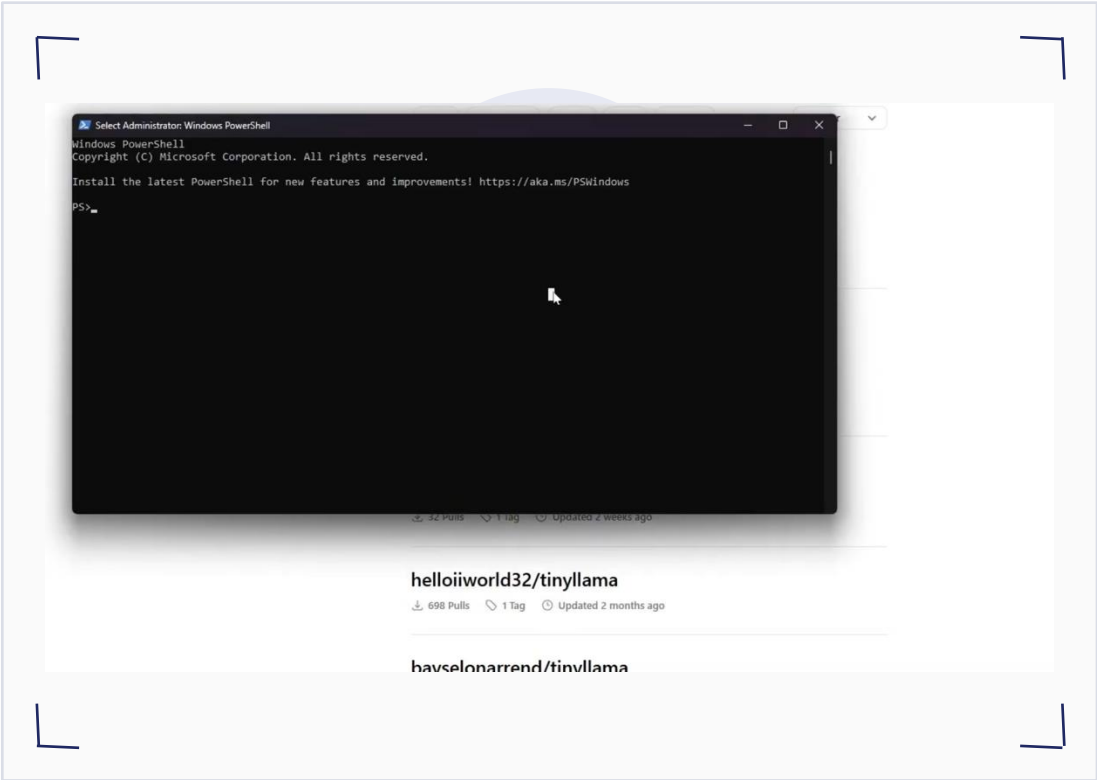
ollama.com — instalador único para Windows, Mac e Linux

- Motor de IA local — baixa e roda modelos abertos (Llama, Qwen, Mistral, DeepSeek) direto no seu computador.
- Todas as próximas peças conectam nele por baixo dos panos.

Perfil de hardware	Modelo sugerido	RAM / VRAM
CPU apenas, sem GPU	phi3:mini	4–8 GB RAM
GPU integrada (Intel/Radeon) — é o meu caso	llama3.1:8b	8–16 GB RAM
GPU dedicada de entrada (6–8 GB VRAM)	llama3.1:8b	16 GB RAM
GPU dedicada 16 GB+ VRAM	gpt-oss:20b	16–24 GB RAM

COMO INSTALAR

```
ollama pull qwen2.5:3b
ollama run qwen2.5:3b
```





AnythingLLM: a interface que eu uso

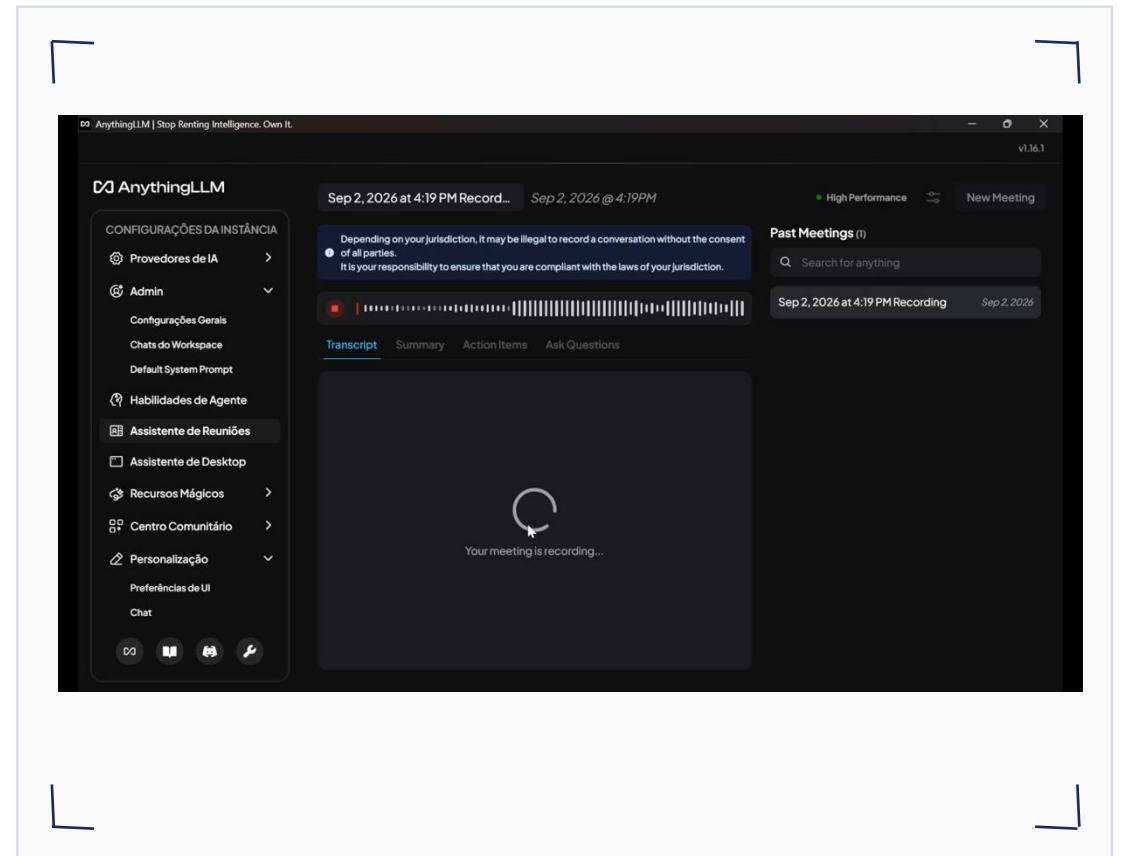
anythingllm.com — instalador único, sem Docker (Windows / Mac / Linux)

- Chat + RAG (documentos e sites) + agente de busca web, tudo pronto de fábrica — zero configuração.
- Fala MCP nativamente e tem "Meeting Assistant": grava, transcreve e resume reuniões sozinho, 100% local (hoje disponível em Mac e Windows).
- Também roda via Docker se você preferir — mais sobre até onde ele chega no "caminho do meio", mais à frente.

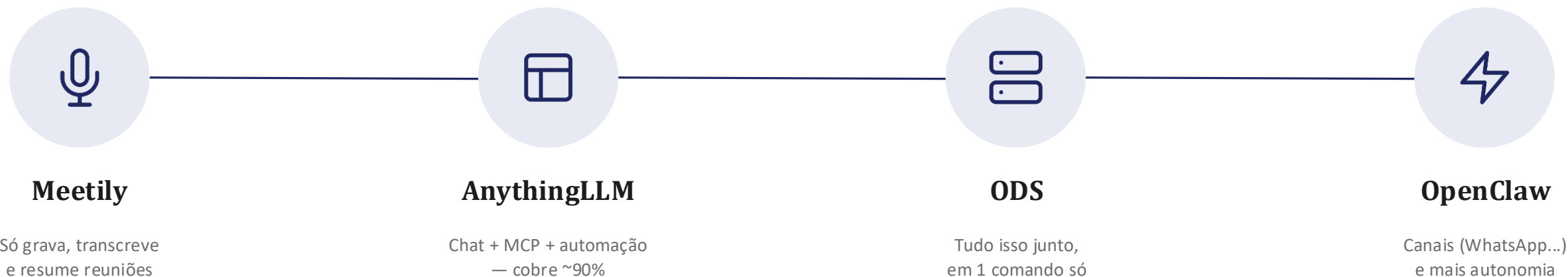
COMO INSTALAR

Instalador único — baixe em anythingllm.com

Direto ao ponto: depois do Ollama (peça 1), esse já é o próximo passo — sem passar por outra interface de chat no meio do caminho.



Do Meetily sozinho ao pacote completo



É a mesma ordem que eu vivi: comecei só com o Meetily, evoluí pro AnythingLLM, descobri que o ODS empacota quase tudo isso sozinho, e testei o OpenClaw quando quis canais de mensagem e mais autonomia.



ODS: o ecossistema inteiro, um comando só

github.com/Osmantic/ODS — Apache 2.0

- Instala e já conecta um motor local, Open WebUI, n8n, busca privada (SearXNG + Perplexica — de onde o Vane nasceu) e geração de imagem (ComfyUI).
- Detecta sua GPU e escolhe o modelo certo sozinho — tem um processo de setup assistido para isso.
- Curiosidade: dentro do próprio ODS, o OpenClaw já aparece como "legacy" — o agente padrão deles agora é um agente próprio, o Hermes Agent.

O que ainda falta

- Sem transcrição de reunião tipo o Meetily
- O Hermes Agent não é um gateway de canais como o OpenClaw

COMO INSTALAR

Linux/Mac: `curl -fsSL install.osmantic.com/ods.sh | bash`
Windows: `rode install.ps1` (ver repositório)





OpenClaw: tudo isso em uma caixa só

openclaw.ai — guia oficial de canais (WhatsApp, Telegram...) em clawdocs.org/guides/channels

WhatsApp

Slack

Telegram

Discord

E-mail

Agenda

Lembretes

Memória

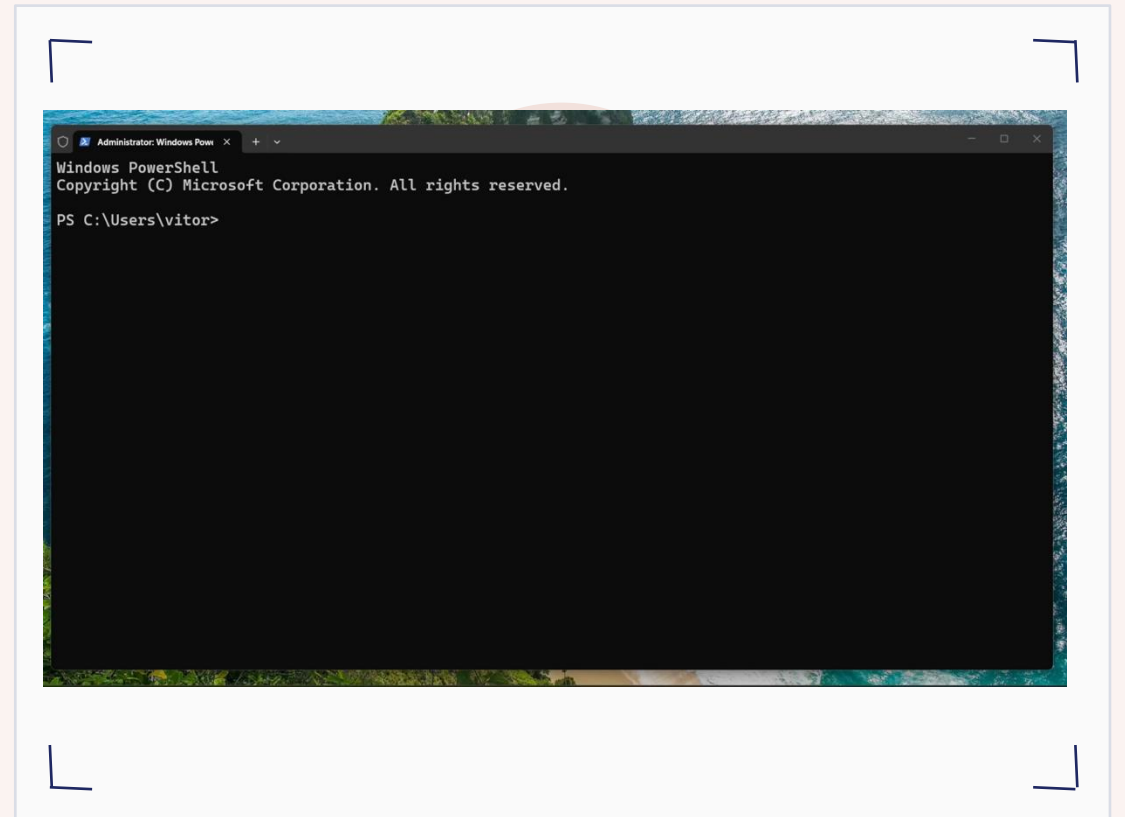


Mais poder, mais responsabilidade

Executa comandos reais na sua máquina — teste em conta/número separado antes de usar no dia a dia.

COMO INSTALAR

```
Windows: iwr -useb https://openclaw.ai/install.ps1 | iex
Mac/Linux: curl -fsSL https://openclaw.ai/install.sh | bash
```





Do DIY ao pacote completo

Peça por peça

- Mais controle sobre cada parte
- Você entende e ajusta tudo
- Mais trabalho de montar e manter
- Ideal pra aprender e pra times

AnythingLLM

- Cobre ~90% do caminho
- MCP + Scheduled Jobs prontos
- Bem menos risco
- Sem canais tipo WhatsApp

ODS

- Tudo isso, empacotado junto
- Detecta GPU e escolhe modelo
- Ainda sem transcrição de reunião
- Ideal pra quem só quer rodar

OpenClaw

- Produtividade quase imediata
- Chega com canais e memória
- Exige mais cautela
- Uso pessoal supervisionado





Todos os links e comandos

- Ollama — ollama.com
- AnythingLLM — anythingllm.com (docs.anythingllm.com)
- Vane (Perplexica) — github.com/ItzCrazyKns/Vane
- n8n — n8n.io
- WhatsApp-MCP — github.com/lharries/whatsapp-mcp
- Meetily — meetily.ai
- Presenton — presenton.ai
- ODS — github.com/Osmantic/ODS
- OpenClaw — openclaw.ai (clawdocs.org)



Aponte para o artigo completo com todos os comandos passo a passo.



Obrigado(a)!

Perguntas?

Vitor Casadei

Senior Data Science Manager, CESAR

linkedin.com/in/vcasadei · vcasadei.com

