

**LINUX
DAY
TORINO
25/10/25**



OPEN SOURCE LLM E COME SFRUTTARLI

GIULIO SCIARAPPA



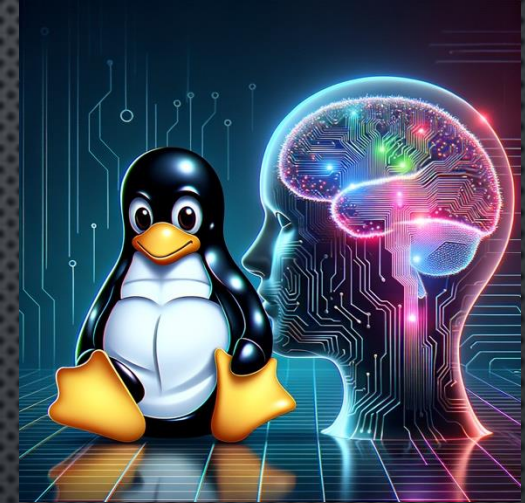
CHI SONO

GIULIO SCIARAPPA

CLOUD ENTERPRISE ARCHITECT



LLM: COSA SONO?



- 🐧 LLM = LARGE LANGUAGE MODEL
- 🐧 ADDESTRATO SU MILIARDI DI TOKEN TESTUALI CON PESI DIFFERENTI
- 🐧 MOTORE DI COMPLETAMENTO DEL LINGUAGGIO NATURALE – «PREVEDONO LA PROSSIMA PAROLA»
- 🐧 BASATO SU TRANSFORMER E SELF-ATTENTION
- 🐧 ESEMPI: CHATGPT, COPILOT, CLAUDE...



LLM FANTASTICI E DOVE TROVARLI



- 🐧 MILIARDI DI PARAMETRI RICHIEDONO CALCOLI COMPLESSI, E GEOMETRICI
- 🐧 LE NORMALI CPU NON SONO OTTIMIZZATE PER QUESTI CALCOLI
- 🐧 SI USANO QUINDI LE GPU, OTTIMIZZATE PER CALCOLI VETTORIALI/GEOMETRICI
- 🐧 PROBLEMA DI MEMORIA -> CARICARE MILIARDI DI PARAMETRI, DOVE? -> RAM (VIDEO)



LLM FANTASTICI E DOVE TROVARLI



🐧 GROSSI DATACENTER E FONDI DI INVESTIMENTO SI SONO FATTI AVANTI E OSPITANO I MODELLI PIÙ RINOMATI

🐧 PRO:

🐧 COSTO ACCESSIBILE

🐧 FACILITÀ DI UTILIZZO

🐧 CONTRO:

🐧 QUANTO SCRIVI È DI LORO PROPRIETÀ

🐧 POTENZIALMENTE FANNO TRAINING DI NUOVI MODELLI CON I DATI DA TE INVIATI



LLM LOCALI, OPEN SOURCE

🐧 LOCALE != OPEN SOURCE

🐧 OPEN SOURCE = SONO NOTI I PESI E IL CODICE CHE LI TRAINA, RENDENDO POSSIBILE IL FINE TUNING

🐧 PROBLEMA-> OSPITARLI LOCALMENTE RICHIEDE POTENZA



PERCHÉ OPEN SOURCE? (E LOCALI...)

- 🐧 AUTONOMIA: NESSUNA DIPENDENZA DA API O CLOUD
- 🐧 PRIVACY: I DATI RESTANO SUL TUO PC
- 🐧 PERSONALIZZAZIONE: PUOI MODIFICARE E FINE-TUNING
- 🐧 COSTO ZERO: NESSUNA SUBSCRIPTION
- 🐧 COMMUNITY DRIVEN: MIGLIORA COSTANTEMENTE



PARAMETRI $\neq$ BIT

PARAMETRI (WEIGHTS)

- 🐧 SONO I VALORI NUMERICI CHE IL MODELLO APPRENDE DURANTE L'ADDESTRAMENTO
- 🐧 OGNI PARAMETRO RAPPRESENTA UNA CONNESSIONE (PESO) TRA NEURONI
- 🐧 UN MODELLO DA **7B** = 7 MILIARDI DI PESI
- 🐧 OGNI PESO È UN NUMERO IN *FLOATING POINT* (ES. FP16 O QUANTIZZATO A 4-8 BIT)
- 🐧

BIT

- 🐧 UNITÀ DI MISURA DELLA MEMORIA (CAPACITÀ DI ARCHIVIAZIONE)
- 🐧 DETERMINANO QUANTO SPAZIO OCCUPANO I PARAMETRI
- 🐧 ES. 7 MILIARDI DI PARAMETRI $\times$ 4 BYTE = **$\sim$ 28 GB DI MEMORIA**



PARAMETRI $\neq$ BIT

Modello	Parametri	Precisione	Peso su disco
Phi-3 Mini	3B	4bit	~ 1,5GB
Mistral 7B	7B	8 bit	~ 7 GB
LLaMA 3 8B	8B	16 bit	~ 16 GB



PARAMETRI $\neq$ BIT

PARAMETRI (WEIGHTS)

- 🐧 SONO I *VALORI NUMERICI* CHE IL MODELLO APPRENDE DURANTE L'ADDESTRAMENTO
- 🐧 OGNI PARAMETRO RAPPRESENTA UNA CONNESSIONE (PESO) TRA NEURONI
- 🐧 UN MODELLO DA **7B** = 7 MILIARDI DI PESI
- 🐧 OGNI PESO È UN NUMERO IN *FLOATING POINT* (ES. FP16 O QUANTIZZATO A 4-8 BIT)

BIT

- 🐧 UNITÀ DI MISURA DELLA MEMORIA (CAPACITÀ DI ARCHIVIAZIONE)
- 🐧 DETERMINANO QUANTO SPAZIO OCCUPANO I PARAMETRI
- 🐧 ES. 7 MILIARDI DI PARAMETRI $\times$ 4 BYTE = **$\sim$ 28 GB DI MEMORIA**





DEMO TIME



MOTORI DI INFERENZA

- 🐧 UN MOTORE DI INFERENZA È IL SOFTWARE CHE CARICA I PESI DEL MODELLO, GESTISCE LA CACHE, ESEGUE IL FORWARD-PASS E RESTITUISCE OUTPUT AL PROMPT.
- 🐧 NON BASTA “AVERE IL MODELLO”: SERVE UN ENGINE OTTIMIZZATO PER LATENZA, THROUGHPUT, MEMORIA E HARDWARE SPECIFICO (CPU/GPU/QUANTIZZAZIONE).
- 🐧 LA SCELTA DEL MOTORE PUÒ FARE LA DIFFERENZA SU **VELOCITÀ** (TIME TO FIRST TOKEN), **USO MEMORIA, SCALABILITÀ**.
- 🐧 IN LOCALE, CON HARDWARE LIMITATO, SONO SPESSO USATI MOTORI LEGGERI/OTTIMIZZATI (ES. LLAMA.CPP) CHE PERMETTONO ESECUZIONE SU CPU.





OLLAMA

- 🐧 *RUNTIME LEGGERO PER MODELLI LLM LOCALI*
- 🐧 *INSTALLAZIONE CON UN SOLO COMANDO*
- 🐧 *GESTIONE DI MODELLI (DOWNLOAD, CACHE, RUN)*
- 🐧 *API REST LOCALI COMPATIBILI CON OPENAI*
- 🐧 *SUPPORTO CPU/GPU (LLAMA.CPP)*



OLLAMA

- 🐧 SFRUTTA QUANTIZZAZIONI GGUF → VERSIONI OTTIMIZZATE PER CPU/GPU CONSUMER;
- 🐧 FORNISCE UN'API REST COMPATIBILE OPENAI, MA LOCALMENTE (PORTA :11434);
- 🐧 GESTISCE CACHE, MODELLI E MEMORIA TRAMITE UN LIVELLO ASTRATTO (NON DEVI COMPILARE NULLA TU);
- 🐧 SU MACOS SFRUTTA METAL / MPS, SU LINUX E WINDOWS CUDA O ROCm (SE DISPONIBILI).



OLLAMA

🐧 *INSTALLAZIONE SU LINUX*

`CURL -fSSL HTTPS://OLLAMA.COM/INSTALL.SH | SH`

🐧 *SU WINDOWS E MAC CON PACKAGE*

[HTTPS://OLLAMA.COM/DOWNLOAD](https://ollama.com/download)

NB: I MODELLI VENGONO SCARICATI IN
~/.OLLAMA/MODELS



OLLAMA

 OLLAMA RUN LLM



DEMO TIME



ALTERNATIVE OSS

GUI

- 🐧 **OPEN WebUI** → INTERFACCIA WEB PER OLLAMA (PIP INSTALL OPEN-WEBUI)
- 🐧 **LM STUDIO** → ALTERNATIVA CROSS-PLATFORM

IDE

- 🐧 PLUGIN VS CODE “CONTINUE”
- 🐧 CHAT AI LOCALE VIA API

DATA LAB

- 🐧 NOTEBOOKS CON LANGCHAIN, LLAMA-INDEX, TRANSFORMERS





DEMO TIME



OTTIMIZZARE LE PERFORMANCE

 QUANTIZZAZIONE (Q4_0, Q5_K_M, ECC.)

 GPU vs CPU FALLBACK

 CACHE DEI TOKEN

 RIDUZIONE CONTEXT WINDOW PER PROMPT
LUNGI

 EVITARE OUTPUT STREAMING SE NON
NECESSARIO



SICUREZZA E PRIVACY

-  NESSUN DATO INVIATO AL CLOUD
-  GESTIONE LOG LOCALE
-  MODELLI VERIFICABILI E AUDITABILI
-  BACKUP DI MODELLI E PESI SU STORAGE PERSONALE



DOVE TROVARE I MODELLI

[HTTPS://OLLAMA.COM/LIBRARY](https://ollama.com/library)

[HTTPS://HUGGINGFACE.CO/MODELS](https://huggingface.co/models)

[HTTPS://GITHUB.COM/OPEN-WEBUI/OPEN-WEBUI](https://github.com/open-webui/open-webui)



RIEPILOGO

- 🐧 CAPITO COSA SONO GLI LLM
- 🐧 VISTI I PRINCIPALI MODELLI OPEN SOURCE
- 🐧 INSTALLATO OLLAMA
- 🐧 ESEGUITO UN MODELLO LOCALE
- 🐧 CREATO IL TUO PRIMO ASSISTENTE PRIVATO





DOMANDE?





GRAZIE!

