

Beyond Bitcoin & FORTH (BBF006)

- SVFIG-2026-07

Title: Extending “word embedding” in Large Language Model to include “stack-words” (stack machine words)

Abstract: We extend “word embedding” in Large Language Model, specifically llama.cpp, to include “stack-words”, which is simply usually “words” in FORTH terminology, with the intention to determine if LLM is capable of generating valid FORTH (and derivatives) expressions.

- FORTH terminology “words”: stack-word, Smurf, Smurd,

1. Use Omnimesg to promote Omnihash and Phoscript?

Hash entered into GitHub page, front end, but JavaScript generate window open link with hash as parameter, then open Phoscript to retrieve message.

- Use this to demonstrate how Omniscientia works!!

2. Omnihash may be used to identify clusters of 10 or N hops nodes with total GPU or CPU and median delay to show where is the best cluster to run free GRATIS tasks.

Discovering cluster computation is the first plus point of GRATIS.

3. Use Omnimesg to send messages amongst GRATIS nodes and establish compute metrics.

4. AGI is overrated if you agree with the following assumptions:

A. The history of human civilization is a process of formulating statistical observations
....

<https://claude.ai/chat/daf1ea2c-7af7-4b85-87c0-d88ecda789d1>

Large language model

67 languages

Article Talk

Read Edit View history

From Wikipedia, the free encyclopedia

"LLM" redirects here. For other uses, see *LLM (disambiguation)*.

A **large language model** (**LLM**) is an **AI model** (typically a **neural network**) trained on a vast amount of text for **natural language processing** tasks, especially **language generation**. LLMs can typically generate, summarize, translate, and analyze text in many contexts, and are a foundational technology behind modern **chatbots**.^[1] Biased or inaccurate training data can make an LLM's output less reliable.^[2]

LLMs are typically based on **transformer** architecture.^[3] **Generative pre-trained transformers** (GPTs) are a type of LLM that is pre-trained to predict the next word.^[4] GPTs are then often **fine-tuned** to follow instructions and to behave as assistants.^[5]

Benchmark evaluations for LLMs attempt to measure **model reasoning**, factual accuracy, **alignment**, and **safety**.^[6]

Part of a series on Machine learning and data mining

Paradigms [show]

Problems [show]

Supervised learning
(classification • regression) [show]

Clustering [show]

Dimensionality reduction [show]

Structured prediction [show]

Anomaly detection [show]

Google

llama.cpp what is embeddings

AI Mode All Images Forums Videos Shopping Short videos More Tools

AI Overview

In *llama.cpp*, **embeddings are numerical vector representations of text that capture its underlying semantic meaning**. Instead of generating new text, the model converts words, sentences, or documents into long arrays of numbers (e.g., [0.085, 0.378, -0.001...]). These numbers allow computer systems to compare different texts mathematically to see how similar they are conceptually. YouTube · Learn Data wi... +4

Core Purpose & Use Cases

While standard LLMs are used for text generation, embedding models are specialized for retrieval and analysis. GitHub +4

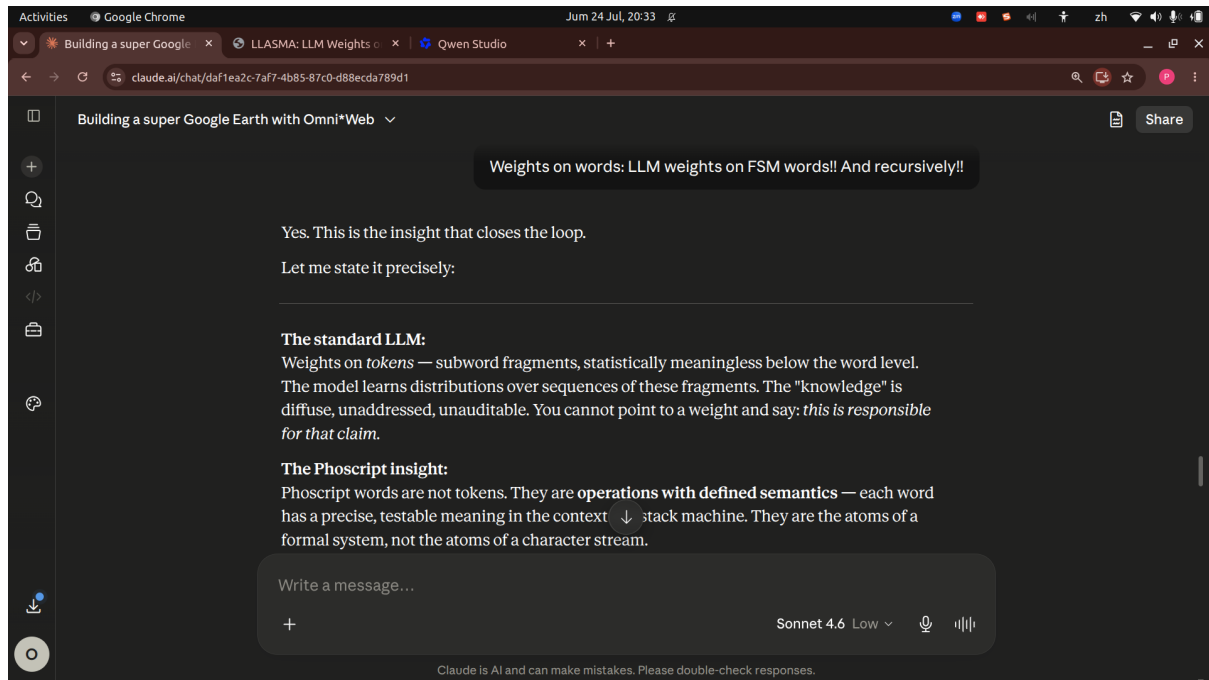
• Retrieval-Augmented Generation (RAG): This is the most common use case. Large

Show more

[illegible]

- Weights on words: LLM weights on FSM words!! And recursively!!
- (A) We started with the prompt as shown above with [Claude.ai](https://claude.ai), which has some previous conversations:

<https://claude.ai/chat/daf1ea2c-7af7-4b85-87c0-d88ecda789d1>

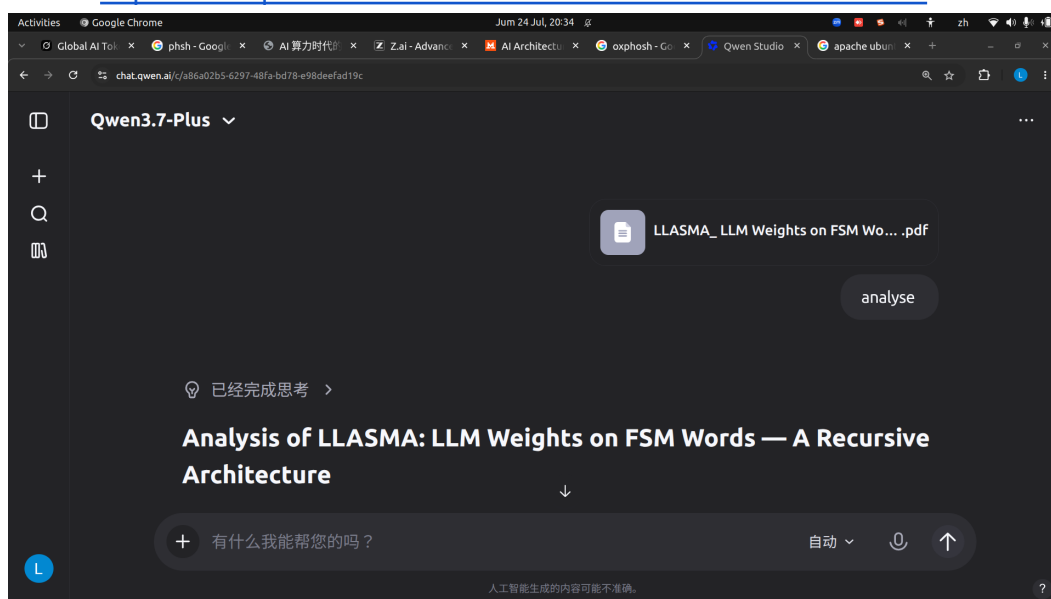


We then save Claude output as PDF, then feed it to 8 other AI services below.

We then fed prompt (B) to all AI services:

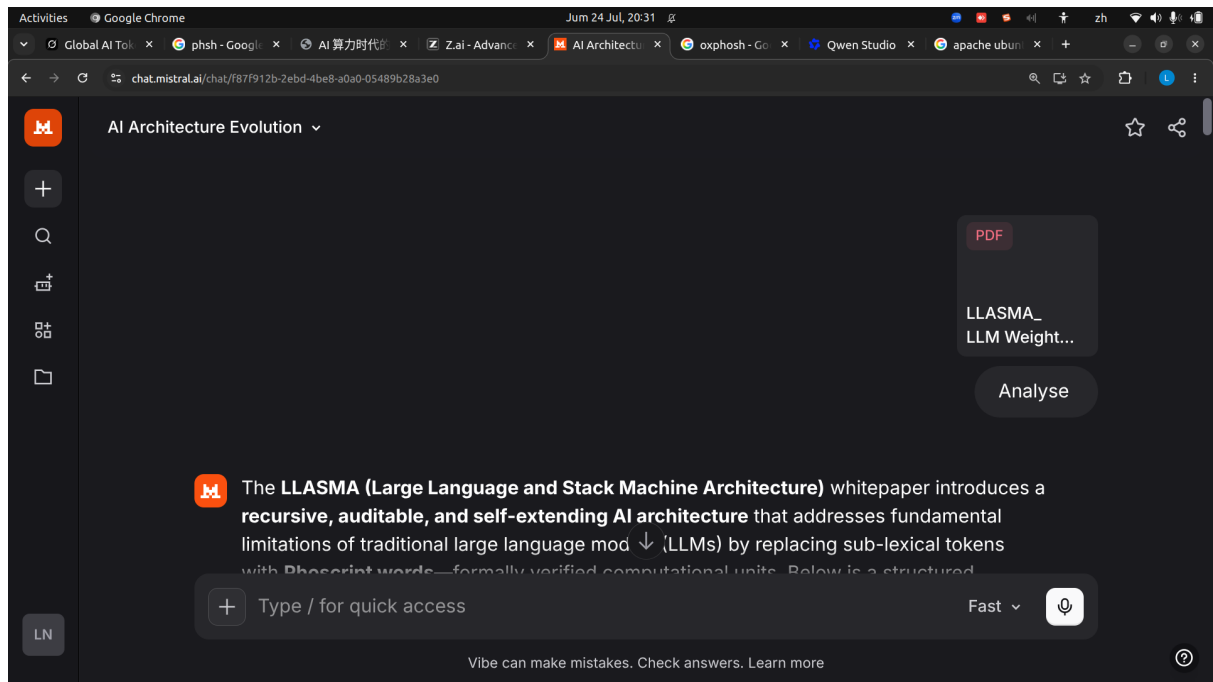
(B) <https://github.com/ggml-org/llama.cpp> describe how we might modify llama.cpp code to predict FORTH stack machine words

- <https://chat.qwen.ai/c/a86a02b5-6297-48fa-bd78-e98deefad19c>



https://chat.qwen.ai/s/t_883beaa6-e5a4-42bb-94a4-77e1947f95c9?fev=0.2.80

3. <https://chat.mistral.ai/chat/f87f912b-2ebd-4be8-a0a0-05489b28a3e0>

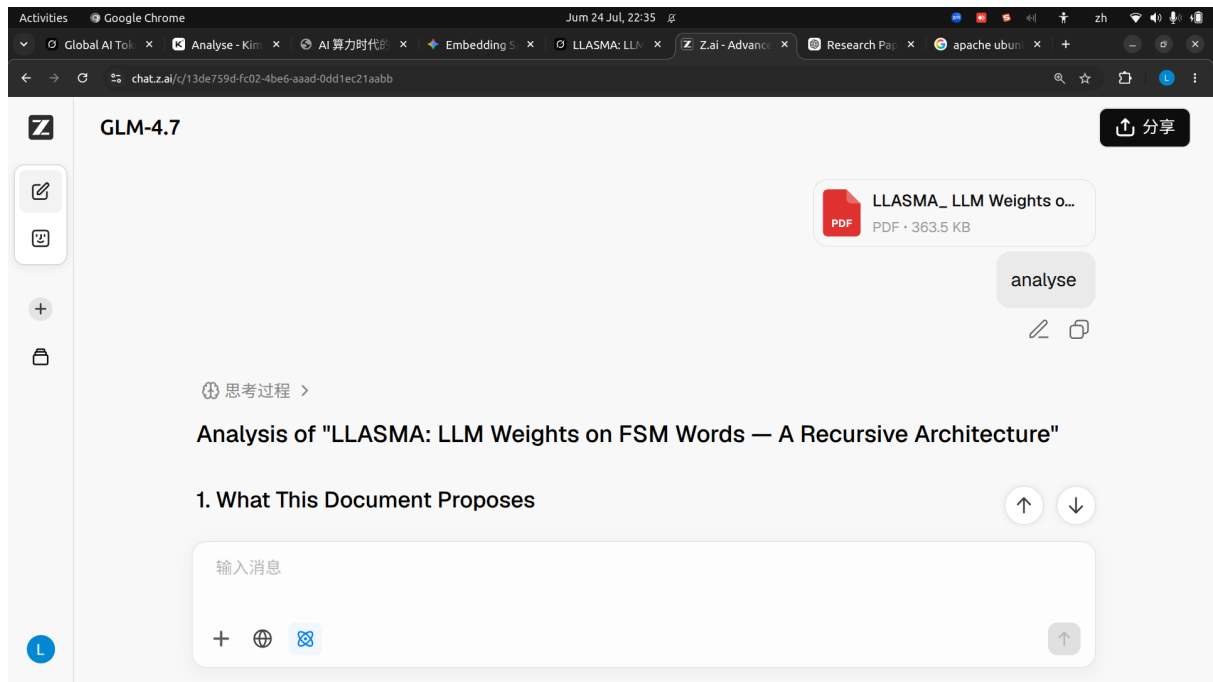


<https://chat.mistral.ai/chat/93d56b6f-07fd-4de9-82e0-50512d210aa9>

4. A. <https://chat.z.ai/c/abdddc2f-a80c-41c8-8809-8f27de0617df>

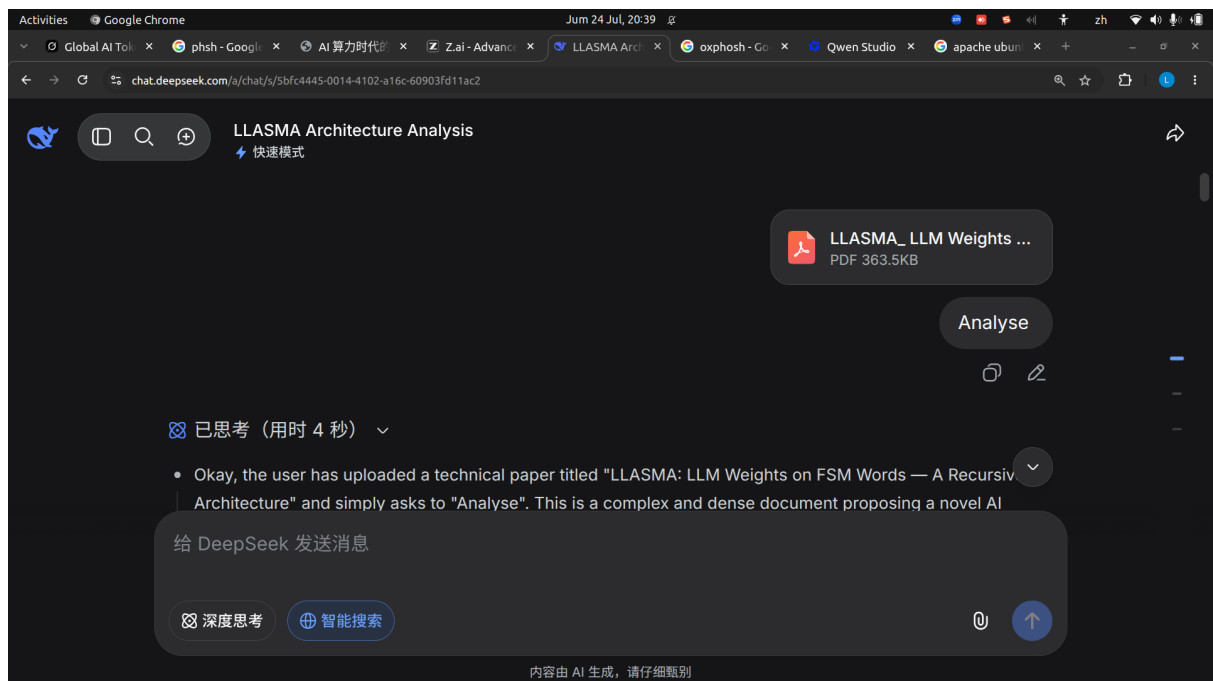


B. <https://chat.z.ai/c/13de759d-fc02-4be6-aaad-0dd1ec21aabb>



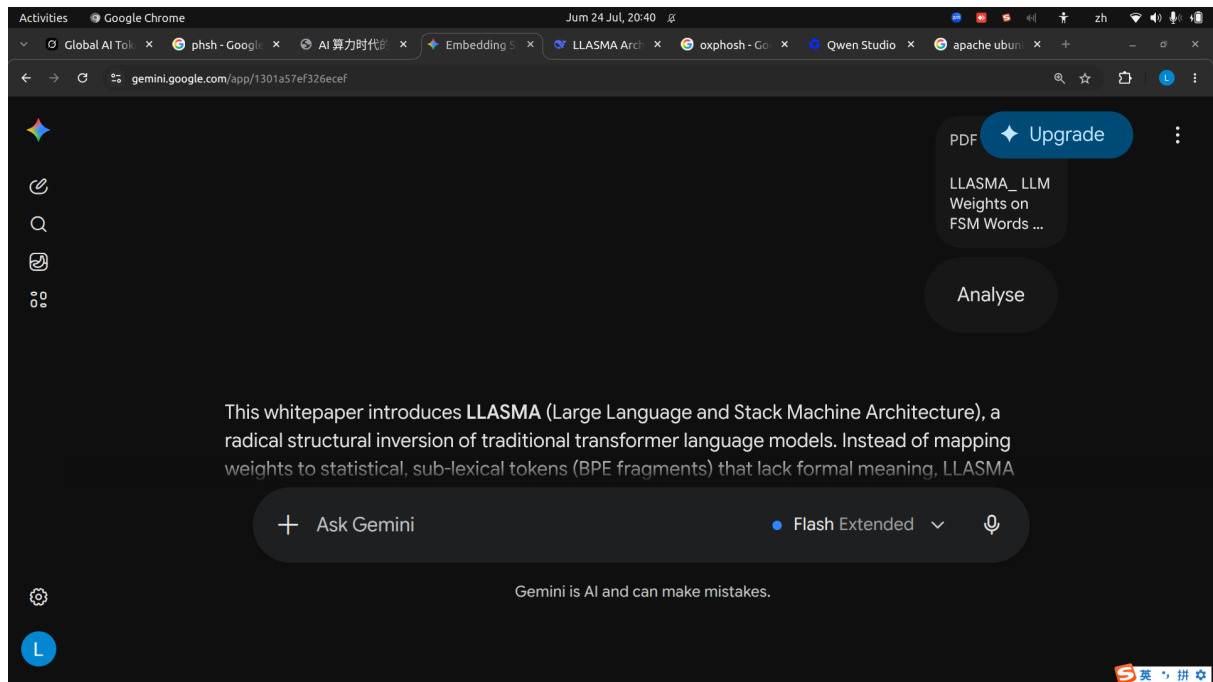
C. <https://chat.z.ai/s/9f857ec3-6522-42ea-930b-fe968384306a>

5. <https://chat.deepseek.com/a/chat/s/5bfc4445-0014-4102-a16c-60903fd11ac2>



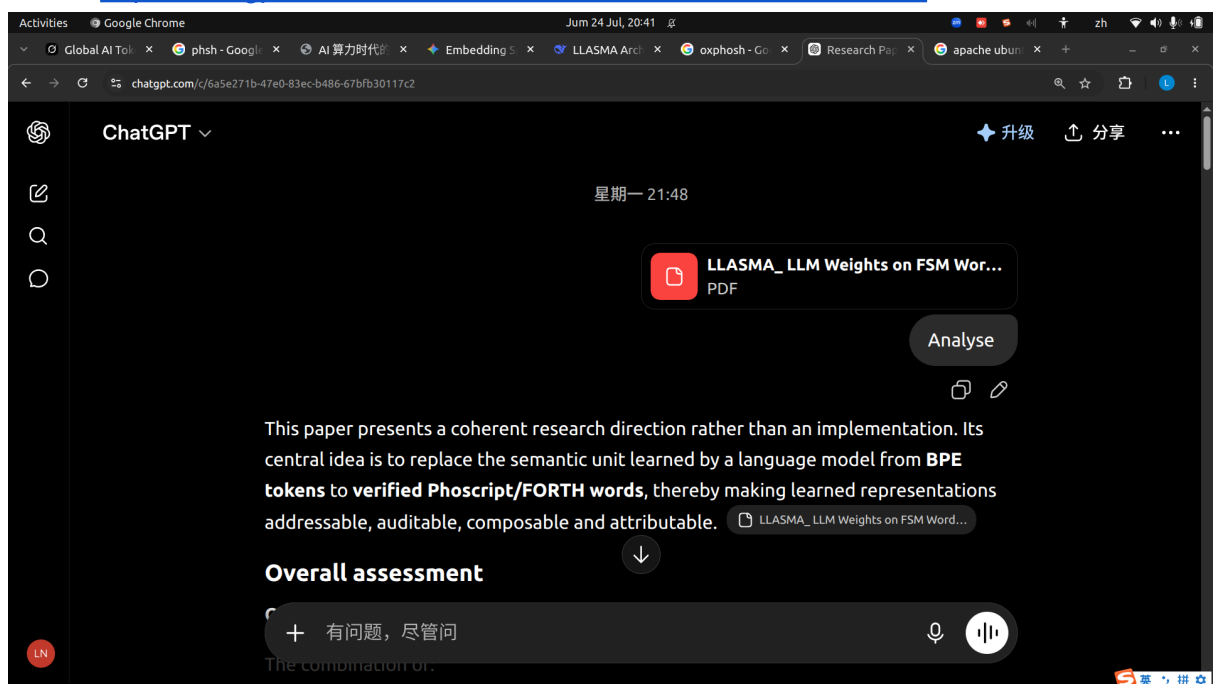
<https://chat.deepseek.com/share/waq5r5cljncepjd0t1>

6. <https://gemini.google.com/app/1301a57ef326ecef>



<https://docs.google.com/document/d/1QnTkVRq0Xhkmo1XgyCAmTyf5oY04al6hm10rkJf8vxl/edit?usp=sharing>

7. <https://chatgpt.com/c/6a5e271b-47e0-83ec-b486-67bfb30117c2>



https://chatgpt.com/s/t_6a63923702488191a6bb9865138efecc

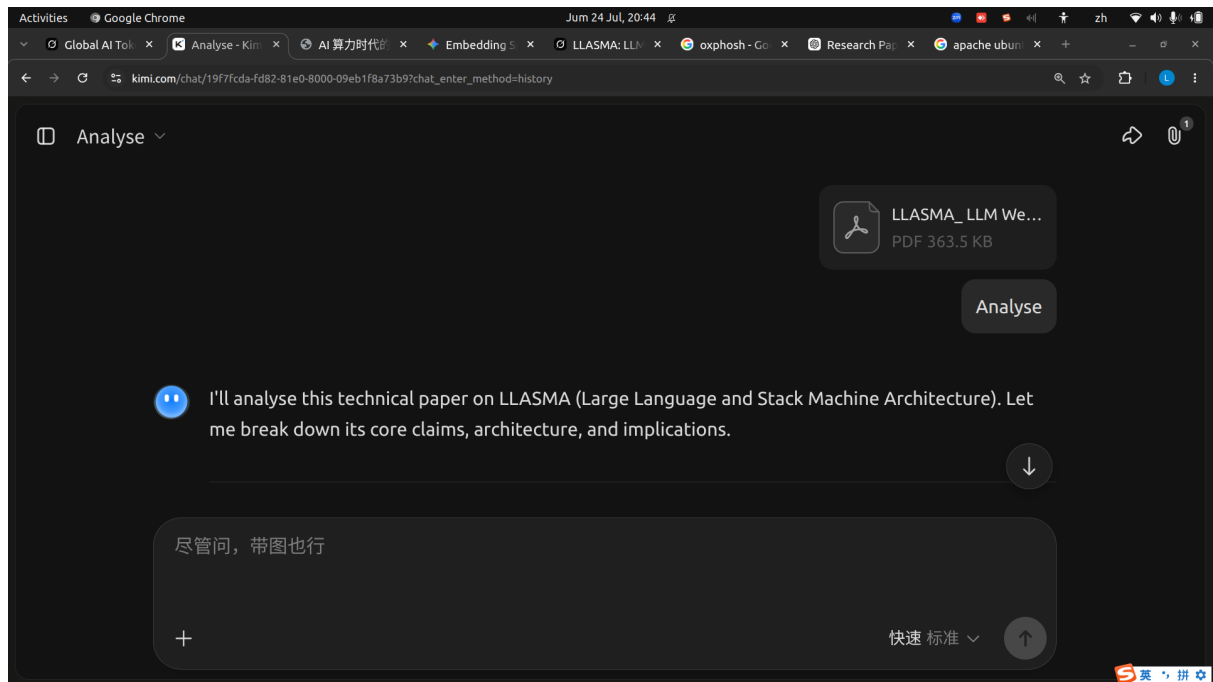
8. <https://grok.com/c/a08b475f-7c16-43de-8d57-332d7708acbc?rid=54ce41f6-f47c-4890-9910-c8212c0972f4>

The screenshot shows the Grok AI web interface in a Google Chrome browser. The address bar displays the URL: <https://grok.com/c/a08b475f-7c16-43de-8d57-332d7708acbc?rid=54ce41f6-f47c-4890-9910-c8212c0972f4>. The interface has a dark theme. At the top, there's a navigation bar with icons for search, edit, share, and other functions. Below this, a prompt area contains a PDF icon and the text "LLASMA_ LLM Weights on FSM Wor...". An "Analyse" button is visible. The main content area shows a response from Grok, starting with "Thought for 6s". The response title is "LLASMA Analysis: LLM Weights on FSM Words — A Recursive Architecture". The text describes the LLASMA (Large Language and Stack Machine Architecture) as a novel AI substrate that replaces the opaque, sub-symbolic token weights of standard transformers with weights over formally defined, verifiable computational primitives in a stack-based language called Phoscript (FORTH-derived). At the bottom, there's a search bar with the placeholder text "+ Ask anything" and a "Fast" dropdown menu.

https://grok.com/share/c2hhcmQtMw_4d37fd2f-0e3f-49d1-877c-bca2e81632a4

The screenshot shows the Grok AI web interface in a Google Chrome browser. The address bar displays the URL: https://grok.com/share/c2hhcmQtMw_4d37fd2f-0e3f-49d1-877c-bca2e81632a4. The interface has a dark theme. At the top, there's a navigation bar with icons for search, edit, share, and other functions. Below this, a prompt area contains a text input: "https://github.com/ggml-org/llama.cpp describe how we might modify llama.cpp code to predict FORTH stack machine words". The main content area shows a response from Grok, starting with "Thought for 19s". The response title is "Modifying llama.cpp for Predicting FORTH/Phoscript Stack Machine Words (LLASMA-style)". The text describes llama.cpp as an excellent base for this because it is highly efficient, modular, and supports custom vocabularies via the GGUF format. The core idea from the LLASMA whitepaper is to replace (or augment) the BPE/subword tokenizer with a Phoscript/FORTH dictionary where each "token" is a full, named word/operation with verified stack effects. Below this, there's a section titled "High-Level Approach" with a list of steps. The first step is "1. Custom Vocabulary: Define Phoscript words as the vocabulary (e.g., +, DUP, SWAP, ELEVATION-BLEND, GPS-READ, etc.).". At the bottom, there's a search bar with the placeholder text "+ Ask anything" and a "Fast" dropdown menu.

9. https://www.kimi.com/chat/19f7fcda-fd82-81e0-8000-09eb1f8a73b9?chat_enter_method=history



点击链接查看和 Kimi 的对话

<https://www.kimi.com/share/19f94eff-b002-8a07-8000-000063253921>