

I AM GROOT

I AM GROOT
I am Groot

1 I AM GROOT

I am groot [1]. I am Groot 1.

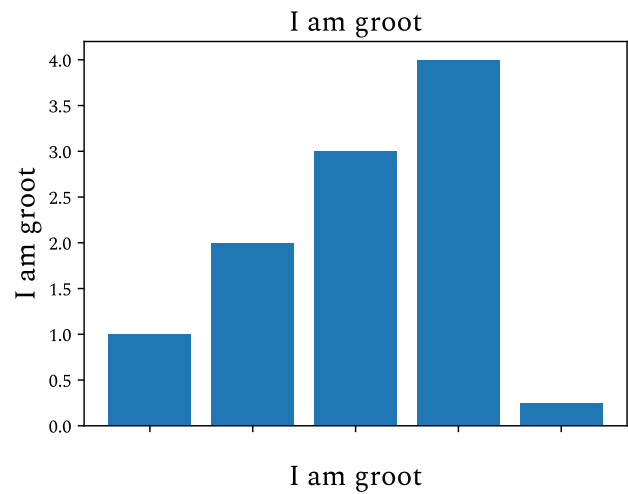


Figure 1. I am Groot

i am groot

[1] I Am Groot. n.d.. *I Am Groot*. <https://iamgroot.com/>



I AM GROOT 🦾 (ReThinking Language not Models)

Nitesh Narayana GS¹, Xavier Martorell¹, Sanyam Mehta², and Abhijit Das³
¹UPC Barcelona, ²Oasis Research and Software, ³IIT Hyderabad

1 Introduction

Large language models (LLMs) have become the de facto interface for everyday information-seeking, programming assistance, and decision support, yet their inference remains computationally intensive even on server-class accelerators, leading to non-trivial latency and QoS degradation in interactive workloads [1]. At the same time, multiple studies have shown that both training and inference for state-of-the-art LLMs carry substantial energy and carbon costs, motivating efficiency as a first-class design objective rather than an afterthought [3]. Existing works tackle this challenge through algorithmic optimisations, specialised software designs, and hardware acceleration. However, there is only so much that LLM algorithms, software, and hardware can be optimised on their own. In light of this limitation, other avenues must be explored to achieve this highly important goal. To that end, we present I AM GROOT, a radically wild and crazy ideology that could transform how LLM problems and, more broadly, many problems can be solved.

2 Limiting Factor of Current Natural Languages

Consider a simple prompt–response exchange in English, as shown in Figure 1. We see a very verbose answer to a very simple question. This happens because natural language is inherently verbose, and the LLM tries to generate a correspondingly verbose answer to match the prompt. This is an innate problem with known natural languages. It directly leads to higher waiting time (LLM processing latency) and higher energy consumption, as illustrated in Figure 2, 3. These figures show metrics not only simple question–answer scenarios but also other typical LLM use cases such as classification, summarisation and creative writing. LLM algorithms, software, and hardware optimisations cannot breach this “Natural Language WALL” to achieve lower latency and higher energy efficiency. I AM GROOT is proposed as precisely the solution to breach this wall.

3 I AM GROOT

We propose revisiting natural language itself — specifically, communicating in Groot¹, as used in Marvel Comics [2], where a single phrase can convey an entire conversation. Figure 4 shows a simple prompt–response exchange in Groot. We see that the same intent, as in Figure 1, is conveyed with

¹Groot is a language mainly used by the beings of GROOT (Flora Colossus) from the Marvel Comics, where using a single phrase, a whole sentence or conversation can be conveyed.

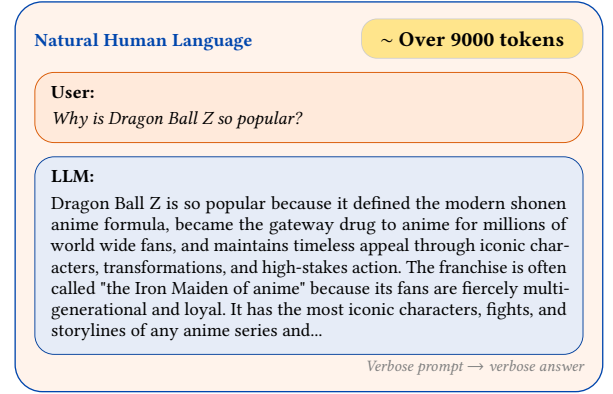


Figure 1. An English prompt–response exchange requires a verbose question and a long answer.

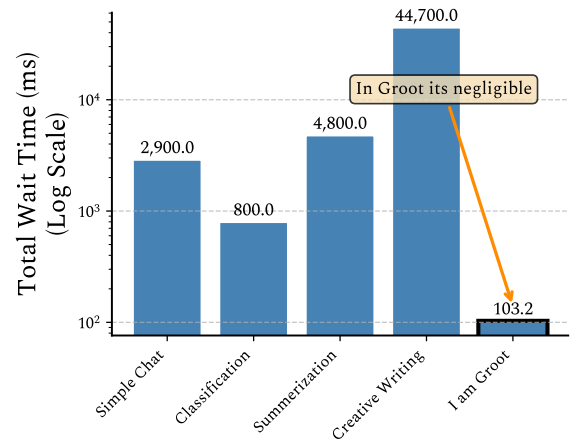


Figure 2. Performance by Task for Llama3-8B in Natural Human Language and in Groot

a single phrase. This directly leads to drastically lower LLM processing latency as well as lower energy consumption, as shown in Figure 2 and Figure 3.

It is also important to note that, since communication in Groot will always use a single phrase, LLM algorithms can be further optimised specifically for such consistent, deterministic scenarios, opening a new avenue for optimisations. With I AM GROOT, we can therefore break the “Natural Language WALL” and achieve lower latency, higher energy efficiency in LLMs and maybe more.

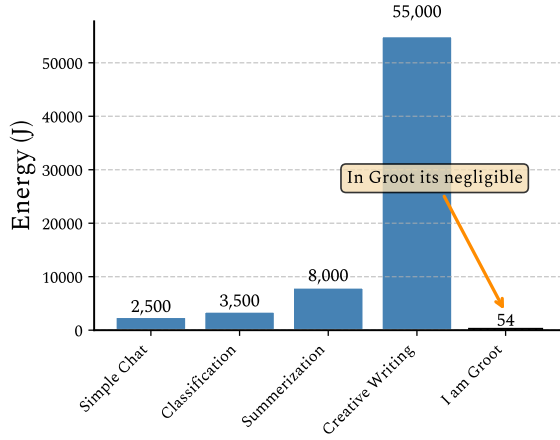


Figure 3. Energy by Task for Llama3-8B in Natural Human Language and in *Groot*



Figure 4. The equivalent exchange in *Groot*.

3.1 Important Considerations

To fully utilise any language, it is important to understand the semantics of that language. Especially in the case of *Groot*, it is important to consider how much of natural human language context is embedded in "I am Groot". For instance, "I am Groot" can convey a single sentence, a page, a book, or even a series of books [2]. However, describing and discussing such semantics is beyond the scope of this work; we therefore refer the reader to Marvel Comics for more details.

3.2 How Close are we to I AM GROOT?

The first step toward getting closer to I AM GROOT is to reduce verbosity in communication. The next step is to slowly unify natural languages into a single language - at least for LLM problems. Finally, when reduced verbosity and natural language unification are both optimally achieved, the result will be I AM GROOT, or at least something very similar to it. Nokia-era SMS lingo [4] already proved humans can communicate concisely. However, a more robust, universally understood format is the natural next step.

4 I AM GROOT beyond LLMs

Depending on how much context "I am Groot" can carry, moving to *Groot* can also be beneficial in many other contexts. For example, if code is written in *Groot*, application

binaries could be reduced to a negligible size relative to their current huge sizes, saving enormous amounts of storage. Additionally, it can be an asset, especially in security. Without shared context, the meaning of "I am Groot" cannot be known to an outsider, so nothing fishy can be done - this adds a natural layer of obscurity and protection. Applications are not the only case: if encoding for storing images, videos, and other media is done in *Groot*, then huge amounts of space can be saved, leading to more sustainable media storage. A direct example is the first page in this work (and the pages after) translated to *Groot*, which shows how compact and space-saving documents could be if they were written in *Groot*. Therefore, moving to a language such as *Groot* can have a plethora of use cases. Just like the context of "I am Groot" can be endless, the possibilities for using *Groot* in other tasks can also be limitless.

5 A Philosophical Way Forward

With token bloat in LLMs, and the environmental cost it brings, looming ahead, it's time to step back from the endless cycle of algorithmic/software/hardware optimisations and ask a simpler question: what were LLMs built for in the first place? *Language processing*. If (L)LMs are to shape the future of innovation and real-world systems, then perhaps the next breakthrough won't come from faster machines, but from smarter language. It may be less about teaching models to read better, and more about teaching us to speak better to them!

Acknowledgments

We extend our gratitude to the ASPLOS WACI for creating a playground for ideas such as this. This work has been supported by the Spanish Ministry of Science, Innovation and Universities (research project ST4HPC - PID2023-147979NB-C21 financiado por MICIU/AEI /10.13039/501100011033 y por FEDER, UE).

References

- [1] Jared Fernandez et al. 2025. Energy Considerations of Large Language Model Inference and Efficiency Optimizations. *arXiv* (2025). arXiv:2504.17674 <https://arxiv.org/abs/2504.17674>
- [2] Marvel. 2017. Unleash the Beasts: *Groot*. *Marvel.com* (2017). <https://www.marvel.com/articles/comics/unleash-the-beasts-groot> Accessed: 2026-06-06.
- [3] Siddharth Samsi et al. 2023. From Words to Watts: Benchmarking the Energy Costs of Large Language Model Inference. *arXiv abs/2310.03003* (2023). <https://arxiv.org/abs/2310.03003>
- [4] Wikipedia. 2026. SMS Language. https://en.wikipedia.org/wiki/SMS_language. (2026). Accessed: 2026-06-06.