

LaVy: Vietnamese Multimodal Large Language Model

Chi Tran

Hanoi University of Science and Technology
chi.tb200083@sis.hust.edu.vn

Huong Le Thanh

Hanoi University of Science and Technology
huonglt@soict.hust.edu.vn

Abstract

Large Language Models (LLMs) and Multimodal Large language models (MLLMs) have taken the world by storm with impressive abilities in complex reasoning and linguistic comprehension. Meanwhile there are plethora of works related to Vietnamese Large Language Models, the lack of high-quality resources in multimodality limits the progress of Vietnamese MLLMs. In this paper, we pioneer in address this by introducing LaVy, a state-of-the-art Vietnamese MLLM, and we also introduce LaVy-Bench benchmark designated for evaluating MLLMs’s understanding on Vietnamese visual language tasks. All code and model weights are public at <https://github.com/baochi0212/LaVy>

we introduce LaVy, Vietnamese first MLLM and achieve state-of-the-art performance in Vietnamese vision language tasks. LaVy is designed to leverage the rich visual and linguistic information present in Vietnamese data, enabling it to tackle a wide range of multimodal tasks with improved performance. Our model outperforms a multilingual baseline mBLIP (Geigle et al., 2023) on different tasks by a large margin. By developing LaVy, we aim to bridge the gap between Vietnamese LLMs and MLLMs, providing researchers and practitioners with a powerful tool for exploring the intersection of language and vision in the Vietnamese context.

1 Introduction

In recent years, Large Language Models (LLMs) have demonstrated remarkable capabilities in various natural language processing tasks, showcasing their proficiency in complex reasoning and linguistic comprehension. The success of LLMs has inspired researchers to explore the potential of Multimodal Large Language Models (MLLMs), which incorporate visual information alongside textual data. MLLMs have shown promising results in tasks that require understanding the interplay between language and vision, such as image captioning, visual question answering, and multimodal machine translation.

While there has been significant progress in developing Vietnamese LLMs, the lack of high-quality multimodal resources has hindered the advancement of Vietnamese MLLMs. The availability of diverse and well-annotated datasets is crucial for training and evaluating MLLMs, as they rely on the integration of visual and textual information to perform multimodal tasks effectively.

To address this limitation and foster research in Vietnamese multimodal language understanding,

Furthermore, to facilitate the evaluation and comparison of Vietnamese MLLMs, we propose the LaVy-Bench benchmark. This benchmark consists an open VQA task and an in-the-wild test set, specifically designed to assess the visual language understanding and generation capabilities of MLLMs in the Vietnamese and in-the-wild images. By establishing a standardized evaluation framework, we aim to promote the development and benchmarking of Vietnamese MLLMs, driving innovation and collaboration within the research community.

In this paper, we present LaVy and the LaVy-Bench benchmark as significant contributions to the field of Vietnamese multimodal language understanding. We provide a detailed description of LaVy’s architecture, data curation and training procedure. Additionally, we introduce the LaVy-Bench benchmark, discussing its design principles, task composition, and evaluation metrics. Through extensive experiments and analysis, we demonstrate the effectiveness of LaVy and the utility of the LaVy-Bench benchmark in advancing Vietnamese MLLM research.

2 Related work

2.1 Large Language Model

Recent advancements in Large Language Models (LLMs) have showcased remarkable capabilities across various natural language processing tasks, including dialogue, creative writing, and problem-solving. Models such as LLaMA (Touvron et al., 2023a,b), Mistral (Jiang et al., 2023), and Gemma (Mesnard et al., 2024) have leveraged scalable Transformer-based architectures (Vaswani et al., 2017) and large-scale data to become foundation models for general reasoning tasks. These models have demonstrated impressive performance and have set new benchmarks in the field.

Following the trend of LLMs, several Vietnamese language models emerged such as PhoGPT (Nguyen et al., 2023b), Vistral (Nguyen et al., 2023a) perform outstandingly in Vietnamese LLM benchmarks and NLP tasks.

2.2 Multimodal Large Language Model

Witness the exceptional performance of GPT-4 (Achiam et al., 2023) and Gemini Pro Vision (Anil et al., 2023) in visual language tasks, recent research has focused on developing Multimodal Large Language Models (MLLMs) to achieve unified understanding and reasoning across different modalities, building upon the success of Large Language Models (LLMs). Various methods have been proposed to integrate information from multiple modalities into pre-trained LLM architectures. For instance, Flamingo (Alayrac et al., 2022) and BLIP-2 (Li et al., 2023) introduce different techniques for fusing visual tokens with frozen LLMs using gated attention or Q-former. Inspired by the effectiveness of instruction tuning, LLaVA (Liu et al., 2023) and MiniGPT-4 (Zhu et al., 2024) align visual input with LLMs through visual instruction tuning, demonstrating impressive results. Another active line of work is researching efficient MLLMs, resulting in lightweight model families like Bunny (He et al., 2024). Meanwhile, recent works are pioneering development in vision-language tasks for low-resource languages, such as Peacock (Alwajih et al., 2024)

3 LaVy

3.1 Architecture

Our model are built with LLaVA architecture (Liu et al., 2023) using three main components:

- **Vision Encoder:** The CLIP-Large model (Radford et al., 2023), is used as the vision encoder.
- **MLP Projector:** A two-layer Multi-Layer Perceptron (MLP) projector is employed to align the output representations from the visual and language modalities. This projector ensures that the visual and textual information is transformed into a common space.
- **Language Model:** The third component is a Large Language Model, which is a language model responsible for generating textual information, take the aligned representations from the MLP projector.

3.2 Data Curation

The barrier for Vietnamese MLLMs' development is resource for training, which is handled by the our novel pipeline of data collection.

- **Translated then Refined:** We utilize LLaVA training data composing of filtered 558K LAION-CC-SBU captions and 150K GPT-generated multimodal instructions. With the acknowledgement of the insufficient competency of open-source translation projects and translation API like VinAI Translate (Nguyen et al., 2022), Google Translate, ... in convert English data to Vietnamese high-quality data, we firstly translate a sample with VinAI translation, then we prompt Gemini Pro to rewrite it in more accurate and natural way for Vietnamese language by pairing translated sample and original English sample in Gemini prompt. Because the captions dataset is massive, we only refine a subset of 150K randomly sampled captions, and combine it with 558K non-write captions to finalize a captioning dataset. Finally, GPT-generated instructions is all refined to make 150K high-quality Vietnamese instructions
- **Synthetic:** Understanding the difference between Vietnamese images and LLaVA's images, we crawl 8.000 images from web in various topics (for example: Vietnamese event images with keyword *Ảnh sự kiện Việt Nam*) and prompt Gemini Pro Vision to generated concise and detailed description of them to enhance LaVy's performance in Vietnamese images. Totally, we have 16.000 Vietnamese

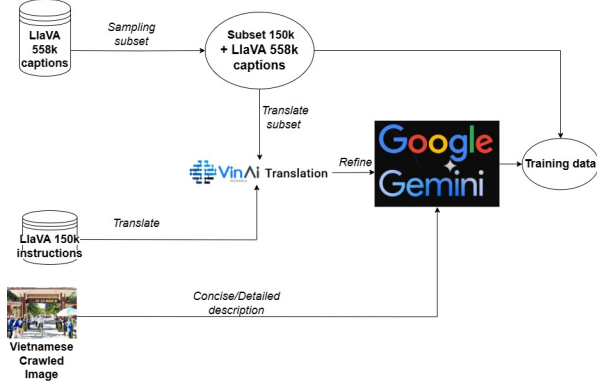


Image 1: Data pipeline

descriptions for crawled images and merge them with rewritten instructions

In eventuality, we curated Vietnamese datasets with 708K image-caption pairs for pretraining and 166K high-quality instructions for finetuning in training step. Our pipeline is clearly illustrated in Image 1.

3.3 Training Procedure

The training procedure is divided into 2 steps:

- **Pretraining:** Align the vision embeddings from a pre-trained vision encoder with the text embeddings from the LLM by optimizing only the cross-modality projector using a cross-entropy loss for next token prediction.
- **Finetuning:** We apply visual instruction tuning to fully utilize the MLLM’s capabilities across different multimodal tasks. We use the same cross-entropy loss as in the pre-training stage, but this time, they employ Low-Rank Adaptation (LoRA) to train both the cross-modality projector and the LLM backbone.

4 Experiment

4.1 Implementation details

We use Vistral 7B as LLM backbone and CLIP large visual encoder. The training process for the LaVy consists of two stages. In the first stage, the model is pretrained using a dataset of 708k captions for 1 epoch, with a global batch size of 64 and a learning rate of 1e-3. During this stage, all model parameters are frozen, except the MLP layers. Besides, we don’t shuffle data but train the model to learn from unrefined data to refined data.

The second stage involves finetuning the model using an instructions dataset. This stage also spans

Model	Accuracy
mBLIP (mT0-XL-5B)	20.0
mBLIP (BLOOMZ-7B)	27.9
LaVy	35.2

Table 1: Zero-shot VQA on OpenViVQA dev set. Models’ output accuracy are evaluated by Gemini Pro

1 epoch, with a global batch size of 32 and a learning rate of 2e-5. In this phase, only the newly introduced LoRA (Low-Rank Adaptation) parameters are trainable.

Besides, in evaluation, we apply greedy decoding to generate all models’ responses (Lin and Chen, 2023)

4.2 LaVy-Bench

We construct LaVy-Bench to benchmark models’ Vietnamese Visual Language understanding.

4.2.1 Zero-shot Visual Question Answering (VQA)

We evaluate the zero-shot Visual Question Answering (VQA) performance of models on the OpenViVQA (Nguyen et al., 2023c) dev set, which consists of 3,505 samples. This dataset challenges the models’ understanding of the relationships between Vietnamese images and natural language. Furthermore, we propose a new metric for automatic evaluation to replace older metrics, such as BLEU (Papineni et al., 2002), which do not accurately reflect models’ competency in the VQA tasks. Our metric is inspired by LLM-as-a-Judge (Zheng et al., 2023), which utilizes Gemini Pro to verify the accuracy of generated responses for question-answer pairs. In Table 1, it’s apparent that LaVy’s zero-shot VQA performance (35.2%) outshadows mBLIP-Bloomz-7B (27.9%) and mBLIP-mT0-XL-5B (20.0%). However, roughly half of the OpenViVQA dataset is composed of TextQA samples (1,733 samples) that do not appear in our training dataset, making OpenViVQA particularly challenging for our model, not mention to our training instructions just includes descriptions of 8.000 Vietnamese crawled images.

4.2.2 In-the-wild benchmark

To further assess the models’ comprehension, we follow the evaluation methodology LLaVA benchmark (in-the-wild) (Liu et al., 2023) and recollect set of 24 diverse images and 60 questions in 3 main types: Complex Reasoning, Detail Description and

Conversation. The collected images and manually crafted questions aim to diversify the test set in various aspects: culture, race, image types,... and avoiding opting for images in crawled training images. For each question, we then prompt Gemini Pro Vision to generate detailed description of image and answer the question, and finally use Gemini Pro to rate the models’ response on the scale of 1 to 5 while taking Gemini Pro’s response as ground-truth. We opted against using detailed descriptions as references (unlike LLaVA) due to their occasional irrelevance. Instead, we found Gemini Pro Vision’s in-the-wild benchmark responses to be closely pertinent. Then we ask model to score outputs in 3 criteria: relevancy, accuracy and naturalness, and sum them all for final score at the end. In comparison with mBLIP baselines in Table 2, LaVy outperforms sharply in all types of questions: Conversation (+36.5%), Detail Description (+59%) and Complex Reasoning (+51%). In overall, our model is scored 61.5 by Gemini Pro. Some qualitative test cases are depicted in Table 3.

Vietnamese MLLMs, enabling complex reasoning and linguistic comprehension in tasks that involve both visual and textual information.

Furthermore, we have presented LaVy-Bench, a comprehensive benchmark designed specifically for evaluating the performance of MLLMs on Vietnamese visual language tasks. This benchmark provides a standardized platform for assessing the capabilities of Vietnamese MLLMs, facilitating the comparison and advancement of these models. Our model also have proved SOTA performance in comparison with mBLIP baselines in test sets of benchmark.

As future work, we plan to expand the capabilities of LaVy by incorporate diverse instructions to entirely handle challenging tasks like Vietnamese OCR, Object Counting. We hope our work will contribute to the advancement of Vietnamese MLLMs’ development.

Model	Conversation	Detail description	Complex reasoning	All
mBLIP (mT0-XL-5B)	47.0	29.0	28.3	34.8
mBLIP (BLOOMZ-7B)	52.6	29.0	44.0	41.9
LaVy	71.8	46.3	66.5	61.5

Table 2: Performance on in-the-wild benchmark

5 Limitations

Our model still have several limitations:

- Although LaVy exhibits deep understanding in Vietnamese visual language tasks but still faces many challenges, for example: TextQA, due to lack of high-quality annotated data for these tasks.
- Moreover, like other MLLMs, our model still suffers from hallucination where it generates irrelevant information, redundant details or misinformation.

6 Conclusion

In this paper, we have introduced LaVy, a pioneering state-of-the-art Vietnamese Multimodal Large Language Model (MLLM) that aims to address the lack of high-quality resources in multimodality for the Vietnamese language. LaVy represents a significant step forward in the development of

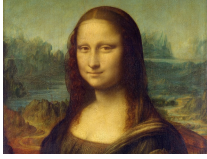
	• Đây là nhân vật lịch sử Việt Nam nổi tiếng nào? Nhân vật lịch sử Việt Nam nổi tiếng trong bức ảnh là Hồ Chí Minh, người sáng lập và lãnh đạo Đảng Cộng sản Việt Nam. Ông là một nhà cách mạng, chính trị gia và nhà lãnh đạo quân sự, đóng vai trò quan trọng trong cuộc đấu tranh giành độc lập của Việt Nam khỏi sự cai trị của thực dân Pháp. Bức ảnh cho thấy Hồ Chí Minh đang phát biểu trước đám đông, thể hiện sự ủng hộ và ảnh hưởng của ông đối với người dân Việt Nam. • Bức ảnh có ý nghĩa như thế nào với lịch sử Việt Nam? Bức ảnh đen trắng này có ý nghĩa lịch sử vì nó khắc họa hình ảnh một người đàn ông Việt Nam, có thể là một nhà lãnh đạo, đang đứng trước một lá cờ đỏ. Lá cờ đỏ tượng trưng cho chủ nghĩa cộng sản và sự đoàn kết của người dân Việt Nam. Bức ảnh có thể được chụp vào thời điểm quan trọng trong lịch sử Việt Nam, khi đất nước đang trải qua những thay đổi chính trị hoặc xã hội. Hình ảnh người đàn ông và lá cờ đỏ thể hiện sự đoàn kết và quyết tâm của người dân Việt Nam trong việc theo đuổi một tương lai tốt đẹp hơn. Bức ảnh này có thể là một lời nhắc nhở về những hy sinh và đấu tranh mà người dân Việt Nam đã trải qua trong quá khứ, cũng như những thành tựu và tiến bộ mà họ đã đạt được.
	• Các con vật nào xuất hiện trong ảnh, và chúng đang làm gì Trong ảnh có một con gấu, một con thỏ và một con mèo. Chúng đang ngồi quanh bàn ăn, cùng nhau thưởng thức bữa ăn. • Viết một câu chuyện có tích cho bức ảnh trên. Ngày xưa ngày xưa, trong một khu rừng xanh tươi, có một chú gấu, một chú mèo và một chú thỏ. Chúng là những người bạn tốt, thường dành thời gian bên nhau, cùng nhau khám phá khu rừng và tận hưởng vẻ đẹp của thiên nhiên. Một ngày nọ, khi đang đi dạo, chúng tình cờ gặp một chiếc bàn nhỏ bằng gỗ, trên đó có một đĩa thức ăn và một tách trà. Chúng rất ngạc nhiên và tò mò, không biết ai đã đặt nó ở đó. Chúng quyết định ngồi xuống và cùng nhau thưởng thức bữa ăn. Khi ăn, chúng trò chuyện vui vẻ, chia sẻ những câu chuyện và kinh nghiệm của mình. Chúng rất vui khi được ở bên nhau và cùng nhau tận hưởng khoảnh khắc này. Sau bữa ăn, chúng tiếp tục cuộc phiêu lưu của mình, khám phá khu rừng và tận hưởng tình bạn của nhau. Và từ đó, chúng trở thành những người bạn tốt nhất, luôn ở bên nhau, bất kể khó khăn hay thử thách nào xảy ra.
	• Bức ảnh được chụp ở sự kiện nào? Bức ảnh được chụp tại lễ trao giải FIFA World Cup, nơi một cầu thủ bóng đá đang cầm chiếc cúp vàng. • Bạn có biết tên của người đàn ông trong bức ảnh là gì? Đúng vậy, người đàn ông trong bức ảnh là Lionel Messi. • Miêu tả cảm xúc của người đàn ông trong bức hình Người đàn ông trong bức hình đang thể hiện cảm xúc vui mừng và tự hào. Anh ta đang cầm chiếc cúp vàng, biểu tượng của chiến thắng và thành tích. Nụ cười của anh ấy cho thấy anh ấy rất hạnh phúc và tự hào về thành tích của mình. Khoảnh khắc này ghi lại niềm vui và sự hài lòng của người đàn ông khi đạt được mục tiêu hoặc thành tựu đáng kể.
	• Bạn có biết ai vẽ bức tranh này không Bức tranh này được vẽ bởi Leonardo da Vinci, một họa sĩ nổi tiếng người Ý thời Phục hưng. Bức tranh nổi tiếng này, có tên là Mona Lisa, được vẽ vào khoảng năm 1503 đến 1506 và hiện đang được trưng bày tại Bảo tàng Louvre ở Paris, Pháp. Bức tranh khắc họa một người phụ nữ với nụ cười bí ẩn, được cho là tự họa hoặc chân dung của Lisa Gherardini, vợ của thương gia Florentine Francesco del Giocondo. Bức tranh được đánh giá cao về kỹ thuật hội họa, sự chú ý đến từng chi tiết và biểu cảm độc đáo của người phụ nữ. • Thảo luận về tác động lịch sử và ý nghĩa của bức tranh này trong thế giới nghệ thuật. Bức tranh Mona Lisa của Leonardo da Vinci là một kiệt tác nghệ thuật có ảnh hưởng lớn đến thế giới nghệ thuật. Bức tranh được vẽ vào thế kỷ 16, thời kỳ Phục hưng, đánh dấu sự phát triển của nghệ thuật và sự ra đời của các kỹ thuật mới như phối cảnh và chiaroscuro. Bức tranh đã trở thành biểu tượng của nghệ thuật Phục hưng và là một trong những tác phẩm nghệ thuật được công nhận và yêu thích nhất trên thế giới. Bức tranh Mona Lisa được biết đến với nụ cười bí ẩn, đã khơi dậy sự tò mò và tranh luận của các nhà sử học, nhà phê bình nghệ thuật và công chúng trong nhiều thế kỷ. Bức tranh cũng đã truyền cảm hứng cho vô số tác phẩm nghệ thuật, văn học và phim ảnh, củng cố vị thế của nó như một biểu tượng văn hóa. Ngoài ra, bức tranh Mona Lisa còn có ý nghĩa lịch sử vì nó được coi là một trong những tác phẩm nghệ thuật có giá trị nhất trên thế giới, với giá trị ước tính lên tới hàng trăm triệu đô la. Bức tranh đã được trưng bày tại nhiều bảo tàng và phòng trưng bày trên thế giới, thu hút hàng triệu du khách mỗi năm, trở thành một phần không thể thiếu trong di sản văn hóa của nhân loại.
	• Bức ảnh phản ánh vấn đề xã hội gì? Có những giải pháp nào để giải quyết vấn đề này? Bức ảnh cho thấy một con phố đông đúc với nhiều người đi bộ và phương tiện giao thông, bao gồm ô tô, xe buýt và xe máy. Tình trạng tắc nghẽn này có thể gây ra nhiều vấn đề xã hội như ô nhiễm không khí, tiếng ồn, chậm trễ giao thông và nguy cơ tai nạn cao hơn. Để giải quyết vấn đề này, có thể áp dụng nhiều giải pháp. Một giải pháp là khuyến khích sử dụng phương tiện giao thông công cộng như xe buýt, giúp giảm số lượng xe cá nhân trên đường. Ngoài ra, thành phố có thể đầu tư vào cơ sở hạ tầng dành cho người đi bộ như vỉa hè và lối đi dành cho người đi bộ để khuyến khích đi bộ và giảm sự phụ thuộc vào phương tiện giao thông. Hơn nữa, các chiến dịch nâng cao nhận thức có thể được triển khai để thúc đẩy các phương thức giao thông thân thiện với môi trường như đi xe đạp hoặc đi chung xe. Cuối cùng, việc thực hiện các quy định giao thông nghiêm ngặt hơn và cải thiện quản lý giao thông có thể giúp giảm ùn tắc và cải thiện chất lượng cuộc sống cho người dân thành phố.

Table 3: Qualitative test cases of LaVy. In-the-wild images and questions are varied in image type, topic, culture...

References

- Achiam, Josh, Adler, Steven, Agarwal, Sandhini, Ahmad, Lama, Akkaya, Ilge, Aleman, Florencia Leoni, Almeida, Diogo, Altschmidt, Janko, Altman, Sam, Anadkat, Shyamal, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Alayrac, Jean-Baptiste, Donahue, Jeff, Luc, Pauline, Miech, Antoine, Barr, Iain, Hasson, Yana, Lenc, Karel, Mensch, Arthur, Millican, Katherine, Reynolds, Malcolm, et al. 2022. Flamingo: a visual language model for few-shot learning. *Advances in Neural Information Processing Systems*, 35:23716–23736.
- Fakhraddin Alwajih, El Moatez Billah Nagoudi, Gagan Bhatia, Abdelrahman Mohamed, and Muhammad Abdul-Mageed. 2024. Peacock: A family of arabic multimodal large language models and benchmarks. *arXiv preprint arXiv:2403.01031*.
- Rohan Anil, Sebastian Borgeaud, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M. Dai, Anja Hauth, Katie Millican, David Silver, Melvin Johnson, Ioannis Antonoglou, Julian Schrittwieser, Amelia Glaese, Jilin Chen, Emily Pitler, et al. 2023. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*.
- Gregor Geigle, Abhay Jain, Radu Timofte, and Goran Glavaš. 2023. mblip: Efficient bootstrapping of multilingual vision-llms. *arXiv preprint arXiv:2307.06930*.
- Muyang He, Yexin Liu, Boya Wu, Jianhao Yuan, Yuezhe Wang, Tiejun Huang, and Bo Zhao. 2024. Efficient multimodal learning from data-centric perspective. *arXiv preprint arXiv:2402.11530*.
- Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, L  lio Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timoth  e Lacroix, and William El Sayed. 2023. Mistral 7b. *arXiv preprint arXiv:2310.06825*.
- Li, Junnan, Li, Dongxu, Savarese, Silvio, Hoi, and Steven. 2023. [BLIP-2: Bootstrapping language-image pre-training with frozen image encoders and large language models](#). In *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 19730–19742. PMLR.
- Yen-Ting Lin and Yun-Nung Chen. 2023. Llm-eval: Unified multi-dimensional automatic evaluation for open-domain conversations with large language models. *arXiv preprint arXiv:2305.13711*.
- Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. 2023. [Visual instruction tuning](#). In *Thirty-seventh Conference on Neural Information Processing Systems*.
- Thomas Mesnard, Cassidy Hardin, Robert Dadashi, Surya Bhupatiraju, Shreya Pathak, Laurent Sifre, Morgane Rivi  re, Mihir Sanjay Kale, Juliette Love, Pouya Tafti, L  onard Hussenot, Aakanksha Chowdhery, Adam Roberts, Aditya Barua, Alex Botev, Alex Castro-Ros, Ambrose Slone, Am  lie H  liou, Andrea Tacchetti, Anna Bulanova, Antonia Paterson, Beth Tsai, Bobak Shahriari, Charline Le Lan, Christopher A. Choquette-Choo, Cl  ment Crepy, Daniel Cer, Daphne Ippolito, David Reid, Elena Buchatskaya, Eric Ni, Eric Noland, Geng Yan, George Tucker, George-Christian Muraru, Grigory Rozhdestvenskiy, Henryk Michalewski, Ian Tenney, Ivan Grishchenko, Jacob Austin, James Keeling, Jane Labanowski, Jean-Baptiste Lespiau, Jeff Stanway, Jenny Brennan, Jeremy Chen, Johan Ferret, Justin Chiu, Justin Mao-Jones, Katherine Lee, Kathy Yu, Katie Millican, Lars Lowe Sjoesund, Lisa Lee, Lucas Dixon, Machel Reid, Maciej Miku  a, Mateo Wirth, Michael Sharman, Nikolai Chinaev, Nithum Thain, Olivier Bachem, Oscar Chang, Oscar Wahltinez, Paige Bailey, Paul Michel, Petko Yotov, Pier Giuseppe Sessa, Rahma Chaabouni, Ramona Comanescu, Reena Jana, Rohan Anil, Ross McIlroy, Ruibo Liu, Ryan Mullins, Samuel L. Smith, Sebastian Borgeaud, Sertan Girgin, Sholto Douglas, Shree Pandya, Siamak Shakeri, Soham De, Ted Klimentko, Tom Hennigan, Vlad Feinberg, Wojciech Stokowiec, Yu hui Chen, Zafarali Ahmed, Zhitao Gong, Tris Warkentin, Ludovic Peran, Minh Giang, Cl  ment Farabet, Oriol Vinyals, Jeff Dean, Koray Kavukcuoglu, Demis Hassabis, Zoubin Ghahramani, Douglas Eck, Joelle Barral, Fernando Pereira, Eli Collins, Armand Joulin, Noah Fiedel, Evan Senter, Alek Andreev, and Kathleen Kenealy. 2024. Gemma: Open models based on gemini research and technology. *arXiv preprint arXiv:2403.08295*.
- Chien Van Nguyen, Thuat Nguyen, Quan Nguyen, Huy Nguyen, Bj  rn Pl  ster, Nam Pham, Huu Nguyen, Patrick Schramowski, and Thien Nguyen. 2023a. [Vistral-7b-chat - towards a state-of-the-art large language model for vietnamese](#).
- Dat Quoc Nguyen, Linh The Nguyen, Chi Tran, Dung Ngoc Nguyen, Dinh Phung, and Hung Bui. 2023b. [Phogpt: Generative pre-training for vietnamese](#). *arXiv preprint arXiv:2311.02945*.
- Nghia Hieu Nguyen, Duong T.D. Vo, Kiet Van Nguyen, and Ngan Luu-Thuy Nguyen. 2023c. [Openvivqa: Task and dataset and multimodal fusion models for visual question answering in vietnamese](#). *arXiv preprint arXiv:2305.04183*.
- Thien Hai Nguyen, Tuan-Duy H. Nguyen, Duy Phung, Duy Tran-Cong Nguyen, Hieu Minh Tran, Manh Luong, Tin Duy Vo, Hung Hai Bui, Dinh Phung, and Dat Quoc Nguyen. 2022. [A Vietnamese-English Neural Machine Translation System](#). In *Proceedings of the 23rd Annual Conference of the International Speech Communication Association: Show and Tell (INTERSPEECH)*.

Kishore Papineni, Salim Roukos, Todd Ward, , and Weijing Zhu. 2002. [Bleu: a method for automatic evaluation of machine translation](#). In *Proceedings of the 40th Annual Meeting of the Associations for Computational Linguistics (ACL)*.

Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. 2023. Learning transferable visual models from natural language supervision. *arXiv preprint arXiv:2103.00020*.

Touvron, Hugo, Lavril, Thibaut, Izacard, Gautier, Martinet, Xavier, Lachaux, Marie-Anne, Lacroix, Timothée, Rozière, Baptiste, Goyal, Naman, Hambro, Eric, Azhar, Faisal, et al. 2023a. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*.

Touvron, Hugo, Martin, Louis, Stone, Kevin, Albert, Peter, Almahairi, Amjad, Babaei, Yasmine, Bashlykov, Nikolay, Batra, Soumya, Bhargava, Prajjwal, Bhosale, Shruti, et al. 2023b. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.

Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In *Advances in Neural Information Processing Systems 30 (NIPS 2017)*.

Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. 2023. Judging llm-as-a-judge with mt-bench and chatbot arena. *arXiv preprint arXiv:2306.05685*.

Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. 2024. [MiniGPT-4: Enhancing vision-language understanding with advanced large language models](#). In *The Twelfth International Conference on Learning Representations*.