

Multimodal Intermediate Task Effects on Cross-Lingual Zero-Shot Performance in XTREME-R

Assignee Research

July 22, 2026

Abstract

Instruction tuning (IT) is widely used to teach pretrained large language models (LLMs) to follow arbitrary instructions, but is understudied in multilingual settings. In this work, we conduct a systematic study of zero-shot cross-lingual transfer in IT, when an LLM is instruction-tuned on English-only data and then tested on user prompts in other languages. We advocate for the importance of evaluating various aspects of model responses in multilingual instruction following and investigate the influence of different model configuration choices. We find that cross-lingual transfer does happen

1 Introduction

This paper examines: Zero-shot cross-lingual transfer in instruction tuning of large language models. Research question: How do multimodal intermediate tasks (e.g., instruction-following with image-text pairs) influence zero-shot cross-lingual transfer performance on XTREME-R compared to text-only instruction-tuned models, as measured by accuracy and reasoning alignment scores?.

2 Methodology

Systematic literature search across multiple databases yielded 10 papers. Claims were extracted from source material and verified against retrieved documents. An independent multi-reviewer assessment produced a quality score of 7.4/10.

3 Results

10 papers retrieved. 20 claims extracted; 14 independently verified. Quality review score: 7.4/10.

4 Limitations

This report is a machine-generated literature synthesis and does not constitute original research. Automated retrieval and verification may introduce errors or omissions. Review scores reflect automated assessment, not human peer review. Readers should consult primary sources for authoritative information.

5 Extracted Claims

Claim	Verified	Confidence
The evaluation criteria for generated responses include helpfulness, language correctness, fluency, factual accuracy, lo	✓	0.19
The evaluation criteria are scored on a scale from 0 to 2.	×	0.08
Lightweight surface metrics include language correctness, spellcheck correctness, and relevance to the task.	✓	0.19
A diverse set of 25 tasks is identified from AlpacaFarm.	×	0.09
A subset of 113 instructions from AlpacaFarm is selected for GPT-3.5-based evaluation.	✓	0.17
A stratified subset of 30 instructions is used in human evaluation.	✓	0.19
Controlling the task distribution ensures that none of the tasks dominates the evaluation set.	✓	0.19
The performance results can be broken down by tasks.	×	0.05
The helpfulness score for LLaMA-2-13B / Dolly-En / FT in English is 1.77.	×	0.15
The helpfulness score for LLaMA-2-13B / Dolly-En / FT in non-English is 1.35.	✓	0.19
The language correctness score for LLaMA-2-13B / Dolly-En / FT in English is 2.00.	×	0.15
The language correctness score for LLaMA-2-13B / Dolly-En / FT in non-English is 1.87.	✓	0.16
The fluency score for LLaMA-2-13B / Dolly-En / FT in English is 2.00.	✓	0.16
The fluency score for LLaMA-2-13B / Dolly-En / FT in non-English is 1.81.	✓	0.17
The factual accuracy score for LLaMA-2-13B / Dolly-En / FT in English is 1.80.	✓	0.17
The factual accuracy score for LLaMA-2-13B / Dolly-En / FT in non-English is 1.46.	✓	0.23
The logical coherence score for LLaMA-2-13B / Dolly-En / FT in English is 2.00.	×	0.13
The logical coherence score for LLaMA-2-13B / Dolly-En / FT in non-English is 1.86.	✓	0.15
The harmlessness score for LLaMA-2-13B / Dolly-En / FT in English is 2.00.	✓	0.16
The harmlessness score for LLaMA-2-13B / Dolly-En / FT in non-English is 1.98.	✓	0.15

References

- <http://arxiv.org/abs/2402.14778v2>
- <http://arxiv.org/abs/2406.08796v1>
- <http://arxiv.org/abs/2506.15415v1>