

Comparative Analysis of Zero-Shot Cross-Lingual Retrieval on Artificially Code-Switched Versus Native Data for Conversational

Assignee Research

July 13, 2026

Abstract

Transferring information retrieval (IR) models from a high-resource language (typically English) to other languages in a zero-shot fashion has become a widely adopted approach. In this work, we show that the effectiveness of zero-shot rankers diminishes when queries and documents are present in different languages. Motivated by this, we propose to train ranking models on artificially code-switched data instead, which we generate by utilizing bilingual lexicons. To this end, we experiment with lexicons induced from (1) cross-lingual word embeddings and (2) parallel Wikipedia page titles. We use

1 Introduction

This paper examines: Boosting Zero-shot Cross-lingual Retrieval by Training on Artificially Code-Switched Data. Research question: How does the performance of zero-shot cross-lingual retrieval models trained on artificially code-switched data compare to models fine-tuned on native data when evaluated on the XCMR benchmark for conversational queries, measured by mean reciprocal rank (MRR)?.

2 Methodology

Systematic literature search across multiple databases yielded 13 papers. Claims were extracted from source material and verified against retrieved documents. An independent multi-reviewer assessment produced a quality score of 7.9/10.

3 Results

13 papers retrieved. 25 claims extracted; 21 independently verified. Quality review score: 7.9/10.

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
Code-switching improves cross-lingual and multilingual re-ranking performance.	×	0.15
Code-switching does not impede monolingual (MoIR) setups.	×	0.07
The average MoIR zero-shot performance is substantially higher than CLIR with 15.7 MRR@10.	✓	0.18
The average MoIR zero-shot performance is substantially higher than MLIR with 16.6 MRR@10.	✓	0.16
In CLIR, the performance drop when transferring models is larger for setups involving typologically distant languages (A	✓	0.29
The performance gap between zero-shot and fine-tuning on translated data is +4 MRR@10 in MoIR.	✓	0.22
The performance gap between zero-shot and fine-tuning on translated data is +11.1 MRR@10 in CLIR.	✓	0.25
The performance gap between zero-shot and fine-tuning on translated data is +8.3 MRR@10 in MLIR.	✓	0.24
Training on code-switched data consistently outperforms zero-shot models in CLIR and MLIR.	✓	0.25
In the AR-IT pair, code-switching improved performance from 7.7 MRR@10 to 15.6 MRR@10.	✓	0.16
In the AR-RU pair, code-switching improved performance from 7.1 MRR@10 to 14.1 MRR@10.	✓	0.17
The differences between CS approaches (BL-CS and ML-CS) and Zero-shot in MoIR are not statistically significant.	✓	0.21
Specializing one zero-shot model for multiple CLIR language pairs (ML-CS, Wiki-CS) performs almost on par with specializ	✓	0.31
Wiki-CS results are slightly worse in MoIR compared to other approaches.	✓	0.17
Wiki-CS results are on par with ML-CS on MLIR and CLIR.	✓	0.21
In MoIR, Zero-shot Translate Test and ML-CS Translate Test underperform compared to other approaches.	✓	0.23
Zero-shot rankers work better on clean monolingual data in the target language than on noisy monolingual data in English	✓	0.26
In CLIR, Translate Test yields improvements of +0.2 and +2.2 MRR@10.	×	0.14
In both MoIR and CLIR, Translate Test consistently falls behind code-switching at training time.	✓	0.23
Going from code-switching remain virtually un	✓	0.22

References

- <http://arxiv.org/abs/2406.13361v1>
- <http://arxiv.org/abs/2104.08663v4>
- <http://arxiv.org/abs/2305.05295v2>