

Impact of SWIM-IR Synthetic Data Scaling on Cross-Lingual Zero-Shot Retrieval Accuracy in XTREME

Assignee Research

July 18, 2026

Abstract

Information retrieval across different languages is an increasingly important challenge in natural language processing. Recent approaches based on multilingual pre-trained language models have achieved remarkable success, yet they often optimize for either monolingual, cross-lingual, or multilingual retrieval performance at the expense of others. This paper proposes a novel hybrid batch training strategy to simultaneously improve zero-shot retrieval performance across monolingual, cross-lingual, and multilingual settings while mitigating language bias. The approach fine-tunes multilingual lang

1 Introduction

This paper examines: Synergistic Approach for Simultaneous Optimization of Monolingual, Cross-lingual, and Multilingual Information Retrieval. Research question: Does increasing synthetic training data volume from 100k to 1 million samples using SWIM-IR improve cross-lingual zero-shot retrieval accuracy on the XTREME benchmark compared to monolingual performance gains?.

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 8.5/10.

3 Results

10 papers retrieved. 13 claims extracted; 12 independently verified. Quality review score: 8.5/10.

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 approach fine-tunes multilingual language models using a mix of monolingual and cross-lingual question-answer pairs	✓	0.41
Experiments on XQuAD-R, MLQA-R, and MIRACL Datasets.	×	0.12
XQuAD-R and MLQA-R are question-answering datasets with parallel questions and passages in 11 languages and 7 languages,	✓	0.20
We report the mean average precision (mAP) for XQuAD-R and MLQA-R.	✓	0.16
The evaluation of the models is conducted on datasets that are completely separate and distinct from the ones used for training	✓	0.23
The results for the Recall metric are in Appendix A.3.1.	✓	0.17
We observe that X-X and X-Y sampling only perform well in monolingual and cross-lingual retrieval settings, respectively	✓	0.24
Our hybrid batch sampling achieves the best performance in multilingual retrieval settings.	✓	0.29
Hybrid batch sampling is better than the other two baseline batch sampling methods when using XLM-R and LaBSE as initial	✓	0.25
Hybrid batch training also substantially reduces language bias in multilingual retrieval compared to monolingual training	✓	0.36
The proposed approach enables strong zero-shot retrieval performance across diverse languages.	✓	0.28
We propose a novel hybrid batch training strategy that simultaneously optimizes retrieval performance across monolingual	✓	0.34
We collect a diverse set of English question-answer datasets and use machine translation to generate parallel question-answer	✓	0.29

References

- <http://arxiv.org/abs/2305.05295v2>
- <http://arxiv.org/abs/2311.05800v2>

- <http://arxiv.org/abs/2408.10536v1>