

Performance Comparison of Monolingual and Multilingual Intermediate-Task Fine-Tuning for Zero-Shot Cross-Lingual Transfer

Assignee Research

July 12, 2026

Abstract

Multilingual contextual embeddings have demonstrated state-of-the-art performance in zero-shot cross-lingual transfer learning, where multilingual BERT is fine-tuned on one source language and evaluated on a different target language. However, published results for mBERT zero-shot accuracy vary as much as 17 points on the MLDoc classification task across four papers. We show that the standard practice of using English dev accuracy for model selection in the zero-shot setting makes it difficult to obtain reproducible results on the MLDoc and XNLI tasks. English dev accuracy is often uncorrelated

1 Introduction

This paper examines: Don't Use English Dev: On the Zero-Shot Cross-Lingual Evaluation of Contextual Embeddings. Research question: How does the performance of monolingual English intermediate-task fine-tuning compare to multilingual intermediate-task fine-tuning (e.g., using XLM-R) on zero-shot cross-lingual transfer for languages with varying script similarity to English, as measured by F1/F1-avg scores on XNLI and TYDI-QA benchmarks?.

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 7.5/10.

3 Results

13 papers retrieved. 10 claims extracted; 8 independently verified. Quality review score: 7.5/10.

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
Published results for mBERT zero-shot accuracy vary by up to 17 points on the MLDoc classification task across four papers	✓	0.34
Selecting checkpoints based on best English dev performance yields zero-shot results varying by up to 15.0% absolute in	✓	0.18
Selecting checkpoints based on best English dev performance yields zero-shot results varying by 4.3% in Thai on the XNLI	×	0.15
During mBERT fine-tuning, English dev accuracy reaches a stable plateau while zero-shot Spanish and Japanese accuracies	✓	0.17
The directional agreement between English dev accuracy and Japanese test accuracy on MLDoc is 42%.	✓	0.17
For all MLDoc languages and 7 of the 14 XNLI languages, variation across 10 runs exceeds 2.5% absolute, while English dev	✓	0.27
A 2.5% difference in accuracy on MLDoc and XNLI test sets is statistically significant at the 5% level using the usual t	✓	0.22
English dev accuracy is often uncorrelated or anti-correlated with target language accuracy in zero-shot cross-lingual e	✓	0.33
Reproducibility issues regarding checkpoint selection and zero-shot performance variation are present in the MLQA task u	×	0.15
Using the target language dev set for checkpoint selection results in dev and test accuracy generally increasing and decreasing	✓	0.17

References

- <http://arxiv.org/abs/2005.13013v2>
- <http://arxiv.org/abs/2005.00633v1>
- <http://arxiv.org/abs/2004.15001v2>