

Impact of Intermediate Language Choice on Cross-Lingual Euphemism Detection Accuracy in XTREME-R

Assignee Research

July 12, 2026

Abstract

Euphemisms are culturally variable and often ambiguous, posing challenges for language models, especially in low-resource settings. This paper investigates how cross-lingual transfer via sequential fine-tuning affects euphemism detection across five languages: English, Spanish, Chinese, Turkish, and Yoruba. We compare sequential fine-tuning with monolingual and simultaneous fine-tuning using XLM-R and mBERT, analyzing how performance is shaped by language pairings, typological features, and pretraining coverage. Results show that sequential fine-tuning with a high-resource L1 improves L2 perfo

1 Introduction

This paper examines: When Does Language Transfer Help? Sequential Fine-Tuning for Cross-Lingual Euphemism Detection. Research question: How does the choice of intermediate language in sequential fine-tuning (e.g., English as a pivot) impact cross-lingual euphemism detection accuracy on XTREME-R compared to direct language pairings?.

2 Methodology

Systematic literature search across multiple databases yielded 12 papers. Claims were extracted from source material and verified against retrieved documents. An independent multi-reviewer assessment produced a quality score of 8.1/10.

3 Results

12 papers retrieved. 6 claims extracted; 5 independently verified. Quality review score: 8.1/10.

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 model is tested on English (EN), Mandarin Chinese (ZH), Spanish (ES), Turkish (TR), and Yorb (YO) after sequential	✓	0.20
The sequential fine-tuning approach involves learning the same task first on one language (L1) and then on a second lang	×	0.13
The number of examples for euphemism (Euph) and non-euphemism (Non-Euph) in the 2025 PETs Datasets are as follows: ZH (2	✓	0.20
The performance of XLM-R and mBERT models on different languages are as follows: XLM-R (EN: 0.821, ES: 0.768, ZH: 0.878,	✓	0.20
The performance of XLM-R and mBERT models on different language pairs are as follows: XLM-R (EN & ES: 0.821, EN & ZH: 0.	✓	0.25
The performance of XLM-R and mBERT models on sequential fine-tuning for different language pairs are as follows: XLM-R (	✓	0.26

References

- <http://arxiv.org/abs/2505.18673v1>
- <http://arxiv.org/abs/2506.15415v1>
- <http://arxiv.org/abs/2508.11831v1>