

Phoneme Error Rate Comparison in Multimodal vs. Speech-Only Self-Supervised Models for Iu Mien Speech

Assignee Research

July 14, 2026

Abstract

The mainstream automatic speech recognition (ASR) technology usually requires hundreds to thousands of hours of annotated speech data. Three approaches to low-resourced ASR are phoneme or sub-word based supervised pre-training, and self-supervised pre-training over multilingual data. The Iu Mien language is the main ethnic language of the Yao ethnic group in China and is low-resourced in the sense that the annotated speech is very limited. With less than 10 hours of transcribed Iu Mien language, this paper investigates and compares the three approaches for Iu Mien speech recognition. Our experi

1 Introduction

This paper examines: Low-Resourced Speech Recognition for Iu Mien Language via Weakly-Supervised Phoneme-based Multilingual Pre-training. Research question: What is the difference in phoneme error rate (PER) between multimodal (speech+text) self-supervised models and speech-only self-supervised models when pre-trained on 1000 hours of data and fine-tuned on 10 hours of Iu Mien speech?.

2 Methodology

Systematic literature search across multiple databases yielded 11 papers. Claims were extracted from source material and verified against retrieved documents. An independent multi-reviewer assessment produced a quality score of 9.0/10.

3 Results

11 papers retrieved. 11 claims extracted; 11 independently verified. Quality review score: 9.0/10.

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
O1 denotes the subword-based monolingual baseline, i.e., trained from scratch.	✓	0.25
M1 denotes the result from fine-tuning of Whistle-small.	✓	0.27
M3 denotes the result from fine-tuning of a Wav2Vec2-base model, trained from scratch using self-supervised pre-training	✓	0.39
Wav2Vec2-cv10 was trained using the fairseq toolkit, following the wav2vec 2.0 base pre-training configuration provided	✓	0.29
M4 denotes the result from fine-tuning of a subword pretraining model, called Mul10-subword, which uses the same pretrain	✓	0.34
Phoneme-based supervised pre-training typically uses International Phonetic Alphabet (IPA) as a common pronunciation ann	✓	0.25
The IPA is designed to provide a unified system of symbols to represent the basic sounds of different languages.	✓	0.23
Each IPA symbol corresponds to a specific phoneme, ensuring a one-to-one relationship between the symbol and the sound.	✓	0.23
With the IPA, it is possible to achieve consistent phonetic transcription across all languages.	✓	0.21
The phoneme-based approach not only allows for efficient model training but also maximizes the sharing of pronunciation	✓	0.27
Our work is based on a recent study on weakly supervised phoneme pretraining, Whistle.	✓	0.21

References

- <http://arxiv.org/abs/2303.06740v1>
- <http://arxiv.org/abs/2407.13292v2>
- <http://arxiv.org/abs/2203.07996v2>