

Multimodal Fine-Tuning Effects on Self-Supervised Speech Model Alignment in Flemish Dutch

Assignee Research

July 19, 2026

Abstract

Recent research in speech processing exhibits a growing interest in unsupervised and self-supervised representation learning from unlabelled data to alleviate the need for large amounts of annotated data. We investigate several popular pre-training methods and apply them to Flemish Dutch. We compare off-the-shelf English pre-trained models to models trained on an increasing amount of Flemish data. We find that the most important factors for positive transfer to downstream speech recognition tasks include a substantial amount of data and a matching pre-training domain. Ideally, we also finetune

1 Introduction

This paper examines: Comparison of Self-Supervised Speech Pre-Training Methods on Flemish Dutch. Research question: What is the effect of multimodal fine-tuning (audio-text) on the alignment capabilities of self-supervised speech models pre-trained on Flemish Dutch, measured via WER on the Common Voice Dutch dataset?.

2 Methodology

Systematic literature search across multiple databases yielded 9 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

9 papers retrieved. 12 claims extracted; 9 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
APC uses a Filterbank feature encoder, GRU aggregator, and reconstructs future frames with an output dimension of 512 and	×	0.14
Mockingjay uses a Filterbank feature encoder, Bidirectional Transformer aggregator, and reconstructs masked frames with	✓	0.16
CPC uses a CNN feature encoder, LSTM aggregator, and identifies future features with an output dimension of 256 and 1.8M	✓	0.17
wav2vec uses a CNN feature encoder, CNN aggregator, and identifies future features with an output dimension of 512 and 3	✓	0.17
wav2vec 2.0 uses a CNN feature encoder, Transformer aggregator, and identifies quantised future features with output dim	✓	0.23
wav2vec 2.0 encoder computes latent speech representations from the raw waveform with 7 temporal convolution blocks.	✓	0.16
wav2vec 2.0 masks a certain proportion of the latent features before feeding to the aggregator.	×	0.14
wav2vec 2.0 uses a quantisation module to map latent feature vectors to discretised versions.	✓	0.15
The final training objective of wav2vec 2.0 is to distinguish the true quantised representation for a masked time step,	✓	0.23
wav2vec 2.0 has two architectures: base with 12 Transformer blocks and large with 24 Transformer blocks in the aggregator	×	0.15
Contextual features at the output of the wav2vec 2.0 aggregator are extracted for downstream tasks.	✓	0.17
wav2vec 2.0 model can be fine-tuned on a labelled set by adding an extra linear layer on top of the context network and	✓	0.15

References

- <http://arxiv.org/abs/2502.17284v1>
- <http://arxiv.org/abs/1805.07467v2>

- <http://arxiv.org/abs/2109.14357v1>