

Spoken document retrieval for an unwritten language

Anonymous ACL submission

Abstract

Speakers of unwritten languages have the potential to benefit from speech-based automatic information retrieval systems. This paper proposes a speech embedding technique that facilitates such a system that we can be used in a zero-shot manner on the target language. After conducting development experiments on several written Indic languages, we evaluate our method on a corpus of Gormati – an unwritten language – that was previously collected in partnership with an agrarian Banjara community in Maharashtra State, India, specifically for the purposes of information retrieval. Our system achieves a Top 5 retrieval rate of 87.9% on this data, giving the hope that it may be useable by unwritten language speakers worldwide.

1 Introduction

Introducing and integrating well-designed digital systems into communities, particularly those with low digital participation, such as oral communities, can enhance their exposure to digital technologies and could reduce inequalities arising from their limited digital use or presence (Deumert, 2014; Gorman et al., 2011).

One application of advancements in language technology is in the application of speech-based search and information retrieval (IR). This task, commonly known as Spoken Document Retrieval (SDR) has been investigated over several decades, mostly notably in DARPA and IARPA programmes such as BOLT, GALE and Babel (Griffitt and Strassel, 2016; Olive et al., 2011; Hartmann et al., 2017). Early work, e.g., (Weintraub, 1993) took the form of simple keyword-spotting tasks (sometimes referred to as Spoken Term Detection), but more sophisticated search capabilities have also been developed (Codem et al., 2002).

In a standard setting, SDR operates over spoken documents (i.e., audio and video files containing

speech) but input queries remain text-based. However, in an alternative setting the input query may be in the form as speech as well. It is this latter formulation that is of course, most relevant to unwritten languages. Whilst SDR systems are typically developed for a specific target language, often using significant quantities of transcribed speech data for model training, this is not possible for an unwritten language. In this case, work to date has adopted a significant simplification of the IR task to that of Query-by-example (QbE), essentially a form of keyword-spotting in which spoken documents are ranked based on the estimated occurrence of an arbitrary spoken input phrase.

QbE systems have been developed in a zero-shot manner (Zhang et al., 2013), meaning that no transcribed data from the target language is required. A simple approach is to perform pattern matching at the acoustic level, usually requiring a variant of dynamic time warping (DTW). However, the advent of unsupervised methods for neural network based acoustic word embedding raises the potential that such embeddings could be used for QbE, or even more sophisticated IR tasks for languages without a written form, or even for languages whose speakers would benefit from voice interfaces but where speech transcription tools are unreliable.

In this paper, we propose a speech embedding technique that we have designed specifically for IR; we find that existing embedding methods such as (Pasad et al., 2024; Sanabria et al., 2023b) are unsuitable for this purpose. We conduct a comprehensive set of development experiments in which we compared the technique to common competing methods – including both DTW and a discrete string search – on a QbE proxy task that we create for several Indic languages. We go on to evaluate the method on a recently-collected corpus of Gormati (Reitmaier et al., 2024), an actual unwritten language. This data was collected specifically with IR in mind, and enables us to evaluate the

performance of our method against metrics that are directly relevant for the community in question. The main contributions of this work are that to our knowledge, this is the first successful approach on a real-world IR task for an unwritten language; and our method is the first to develop speech embeddings for a new language in a zero-resource setting targeted specifically to SDR.

2 Prior work

2.1 Spoken Document Retrieval

The key challenges of SDR are how to represent spoken queries and documents, and how to perform search using those representations. Most approaches use large vocabulary continuous speech recognition (LVCSR) to transcribe both queries and documents into text, followed by standard text IR methods (Chelba et al., 2008). In cases where high recall is required, lattices containing alternative candidate transcriptions can be used in place of a single 1-best transcription (James and Young, 1994; Richardson et al., 1995).

Historically, SDR was performed without the need for word-based transcription by using phonetic transcriptions instead (Amir et al., 2001; Ng and Zue, 2000). This approach was commonly termed “phonetic search”. When both queries and documents are transcribed into sequences of discrete phoneme-like symbols, it is possible to use string-matching algorithms to perform keyword search. However, the matching must be robust to the high error rates typically seen in phone recognition, requiring methods such as Buzo et al.’s (2013) windowed string search method – which calculates string distances between queries and segments of documents – or retrieval with the vector space model (VSM), using phone n-grams as terms (Moreau et al., 2004). It should be noted that phonetic search methods often exhibit very high false positive rates.

When performing QbE or another form of fully speech-speech retrieval, it is also possible to perform matching with continuous representations in the acoustic domain. In this case, it is essential to perform dynamic time warping to account for the differing term lengths between query and document audio. Early work used standard signal processing features such as mel-frequency cepstral coefficients (MFCCs) (Park and Glass, 2008), but such features are not robust to variation in speaker characteristics or acoustic environment (Sudhakar et al., 2023).

Subsequently many alternative neural-network features have been tested, including phone posteriorgrams (Hazen et al., 2009) and multilingual bottleneck features (BNFs) (van der Westhuizen et al., 2022). San et al. (2021) found that self-supervised features from wav2vec 2.0 and XLSR-53 can outperform MFCCs and BNFs using DTW.

For low-resource languages, LVCSR systems may suffer from unacceptably high error rates, or may not be available at all; and of course, for unwritten languages it simply may not be possible to produce word-like output. In such cases, it may be necessary to use phonetic search methods or acoustic domain matching. We compare both of these approaches in our experiments.

2.2 Acoustic word embeddings

Acoustic Word Embeddings (AWEs) are embeddings of speech that aim to capture word-like properties. In theory, they may be able to use contextual information to learn semantic information, in a manner to text-based word embeddings. Compared to text, however, speech data has a much higher time resolution; contains additional nuisance factors that are unrelated to word identity; and is generally available in more limited quantities. Furthermore, word boundaries are generally unknown. However, since they can be trained in an unsupervised manner on a target language – or trained on related languages – AWEs can be useful for untranscribed languages (Sanabria et al., 2023a; Jacobs and Kamper, 2021).

The extent to which AWEs are able to capture semantic information is still a current research topic. Pasad et al. (2024) demonstrate that self-supervised representations (e.g., HuBERT vectors) do contain some level of semantic information, useful for discriminating words. They additionally show that when these features are used as inputs to downstream models, they perform much better at word discrimination than with other more standard features - e.g., MFCCs.

Pasad et al. (2024) show that pooling self-supervised representations (e.g., from wav2vec2 or HuBERT) can produce effective AWEs, and Sanabria et al. (2023b) find HuBERT to be the best for English word discrimination. Although, since HuBERT is only trained on English, the quality degrades when it is applied to other languages. However, the recent release of mHuBERT, a compact model with the same architecture as HuBERT-Base, trained on 147 languages, could enable generating

high-quality AWEs for languages beyond English (Boito et al., 2024; Hsu et al., 2021).

Instead of pooling, one can train a model that uses self-supervised representations to produce AWEs. Sanabria et al. (2023a) describe a simple method of producing AWEs using a learned pooling layer trained on at least one hour of target language speech. They use a multilingual phone recogniser (MPR) to transcribe the recordings and then train the model contrastively to embed speech segments with the same transcription similarly. The limitations of this method are that it is not clear how to apply it to a QbE task and it requires at least one hour of target language training data plus an MPR. In contrast, Hu et al. (2021) and Jacobs and Kamper (2021) explored the performance of transfer learning with AWE models on a QbE task by training on well-resourced languages and applying the models to low-resource target languages without finetuning. Both studies embedded the entire query, segmented the search collection with a sliding window and embedded these segments. Jacobs and Kamper (2021) found that training using languages that are closely related to the target language improves performance, and when adding training languages, the largest improvement is gained from adding a single related language. We hypothesise that adopting this approach with a learned pooling model, with its lower data requirements, could allow for the development of an effective QbE system using an AWE model trained with only a small amount of related language training data.

3 Data

3.1 Gormati Dataset

Our primary task was IR for Gormati, an unwritten language spoken by the Banjara farming community in India. This dataset was recently collected by Reitmaier et al. (2024). Community members were asked to provide a natural spoken description of images of various crops. The dataset contains 302 recordings (3.8 hours) split over 32 different classes/images.

To select Gormati queries, we removed silent recordings and those longer than 3 minutes to avoid memory issues. Any classes with only 1 recording were removed from the corpus. The remaining recordings were divided into queries and documents. We used 99 queries as Reitmaier et al. (2024). The queries were unmodified recordings randomly selected from a given class, and the num-

Language	# Queries	Corpus Size
Gormati	99	288
Gujarati	896	23,255
Hindi	163	4,686
Marathi	89	2,550
Odia	30	873
Tamil	983	28,321
Telugu	984	28,504

Table 1: Data corpus sizes and number of queries.

ber of queries in each class was proportional to the number of recordings in that class. We ensured that classes with only 2 recordings had at least 1 query. This left 288 documents and 99 queries consisting of natural spoken descriptions. Since these descriptions might discuss a topic indirectly rather than directly naming the subject, there is no assurance of any lexical or phonetic overlap between a query and its corresponding documents.

Our processed data had small discrepancies with the data described in Reitmaier et al. (2024). Despite our best efforts we could not reconcile these discrepancies. However, their search collection was restricted to only include high volume classes, while ours has no such restriction, including classes with as few as two recordings. Hence, our formulation should be more difficult and realistic.

3.2 Indic Datasets

Given the limited Gormati data, we used higher-resource Indic language data during the development of our models. We used data for Gujarati, Hindi, Marathi, Odia, Tamil, and Telugu from the 2021 Interspeech Multilingual and Code-Switching (MUCS) challenge (Diwan et al., 2021). For each language, we combined the training and test sets to form the search corpus, we filtered out short utterances under 4 words, and we sampled one example of each repeated sentence.

For QbE, we extracted single-word queries using tf-idf weighting. We imposed a minimum document frequency of 2 and a maximum of 6. We ranked each word by its maximum tf-idf value and selected the top scorers as queries, such that the ratio of queries to corpus size was 0.03-0.04, as in Table 1. For each keyword, we selected one recording as the query source, while the remaining recordings containing the keyword were the corresponding gold standard matches.

4 Methods

4.1 DTW Baseline

Our baseline was motivated by San et al.’s (2021) use of self-supervised features in low-resource DTW-based QbE. For each query, we ranked relevant recordings based on normalised subsequence DTW (Giorgino, 2009; Tormene et al., 2009) with Euclidean local distance over mHuBERT-147 representations (3rd iteration, final layer) (Boito et al., 2024).¹ We extracted query representations by embedding the entire recording, and using gold standard timings to clip the vectors corresponding to the query. This method allowed for better query representations than clipping queries first, which was problematic because of issues related to short audio length and discontinuities at the audio edges.

4.2 Phone Recognition Baseline

We also performed baseline experiments with phone-based ASR. We used the MPR from Reitmaier et al. (2024) to transcribe both queries and documents, then performed string-based search on these transcriptions. We experimented with phone transcriptions from Allosaurus (Li et al., 2020), but found that these led to poor retrieval results. We also tried using discrete mHuBERT codes instead of phone sequences, but found that there were too many distinct codes for retrieval to be effective.

As in Moreau et al. (2004), we used vector-space model (VSM) retrieval on these phone transcriptions. Queries and documents were represented as vectors of tf-idf weighted terms and scored based on their cosine similarity. We used all phone n-grams from 1-grams to 8-grams as terms.

We also implemented a simple approximate string search method, similar to Buzo et al. (2013). For each query, we slid a window over each document transcription, calculating the edit distance between the window and query phones. Documents were given the score of the window with the smallest edit distance, and documents with the lowest scores were returned as best matches. We used a window size of 1.2 times the query length.

4.3 mHuBERT model

We used mHuBERT model to convert documents and queries recordings into a sequence of vectors. Then, we compute cosine similarity between all

¹We piloted DTW experiments with MFCCs, but found that they were unable to cope with speaker variation, leading to poor results.

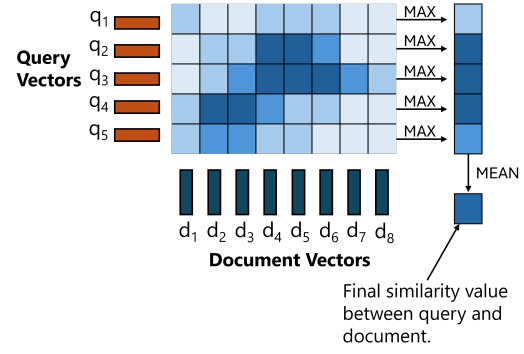


Figure 1: Example of calculating the final similarity value between a query and a document from the cosine similarity matrix. Darker cells indicate higher similarity.

extracted vectors for a given query and document. Finally, for each query vector, the maximum similarity over the document vectors is taken and the similarities are averaged over the query vectors to get a single similarity value between a query and a document, illustrated by Figure 1. This is done for all combinations of queries and documents and for each query, the documents are ranked based on their similarity scores.

4.4 AWE model

We obtained AWEs by learning a pooling function over mHuBERT features. The architecture of the pooling function is the same as in Sanabria et al. (2023a) and Algayres et al. (2022), with a layer norm followed by a 1D convolution, then by a transformer layer with positional embeddings, and finally by a max pooling layer through time (total: 6.8M params). The pooling function is trained using the NTXent contrastive loss as in Sanabria et al. (2023a). As input, this loss takes a batch consisting of several pairs, each from a different class. Within each pair, the two examples serve as positive examples for each other, while examples from other pairs act as negatives, and vice versa for the other pairs. Samples are selected based on their phonetic transcriptions – segments that share the same transcription are considered positive samples. We train using gold phone transcriptions as the training language can be higher resource than the target language and so may have gold transcriptions. When gold labels are unavailable, we experiment with MPR transcriptions in Section 5.6.

We train three models separately on Tamil, Telugu and Gujarati.² During training, we test the

²We choose these languages since they are sampled at 16 kHz, the required sample rate for mHuBERT.

multilingual search performance of models at each epoch by testing on a Marathi search task. We early stop when performance does not improve for 2 epochs. We use the Adam optimiser with learning rate, $l = 10^{-4}$ and we set the NTXent temperature $\tau = 0.15$. Training takes under 5 hours on an NVIDIA V100 16GB (Volta).

We present two inference methods for SDR with AWEs based on a sliding window approach, similar to those discussed in Section 2.2. However, here we window both the document and the query because for Gormati, the queries can be just as long or longer than many of the documents.

The first method, *Phone Window Inference* (Figures 2 and 4), relies on the phone timings from MPR to divide recordings into segments of continuous, non-silent phones.³ The min/max length of these segments is specified in phones. For example, for 2-4 phones, all segments containing 2 continuous, non-silent phones are extracted first, followed by extracting segments with 3, and 4 phones. These segments are then embedded using the AWE model, and queries and documents are compared using the cosine distance method, described in Section 4.3.

The second method, *Time Window Inference* (Figures 3 and 4), does not require phone timing knowledge. Instead, an average phone length of 80 ms is assumed and the window is applied as a standard sliding window with 50% overlap. In this case, 2-4 phones would correspond to windows of 160 ms, 240 ms, and 320 ms.

We anticipate that phone window inference with gold labels will outperform time window inference and phone window inference with an MPR, because of the additional noise from silences and partial phones in the time window and from incorrect transcriptions with the MPR. However, as we are unlikely to have gold labels for inference we treat it as a top line system.

4.5 Evaluation Metrics

Reitmaier et al. (2024) evaluated their Gormati voice search system with a *Top 5* metric, which is the percentage of queries that had at least one correct document in their top 5 returns. They used this metric because their voice search app displayed 5 images per page and users could reliably identify a single correct image among them. We use this metric, not just for consistency but also because our

³Note that, for training and inference, recordings are first passed through mHuBERT before they are split up into different segments.

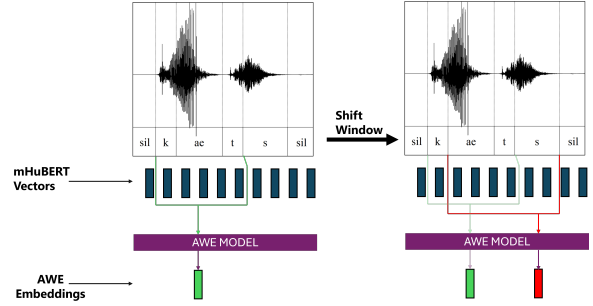


Figure 2: Example phone window inference with a length of 3.

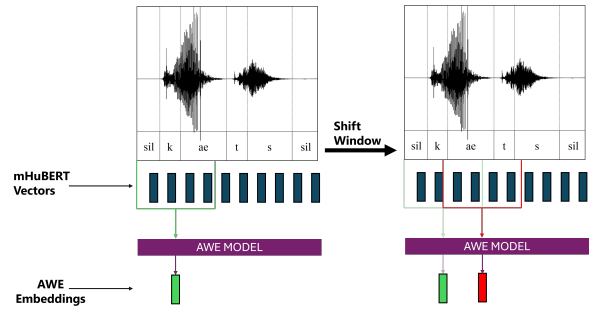


Figure 3: Example of time window inference.

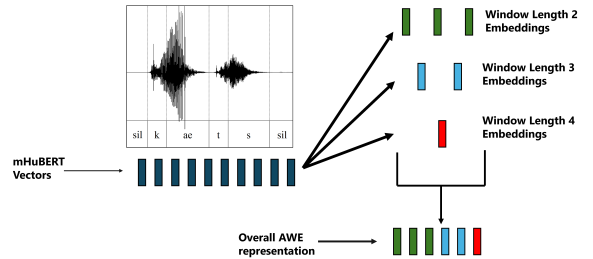


Figure 4: Example AWE representation for a recording using either phone or time window inference with a length of 2-4.

systems are designed for Gormati. Consequently, they could be integrated into the existing app, and should be evaluated accordingly. However, the Top 5 metric is quite coarse - it does not consider the number of results in the top 5 or their order. Likewise, it does not indicate how the system performs across all returns. Hence, in addition to the Top 5 metric, we use Mean Average Precision (MAP) and Mean Average Precision at 5 (MAP@5).

5 Results

5.1 MPR Quality

We first evaluated the quality of the MPR on Tamil. We transcribed all Tamil queries with the MPR and treated these as “reference” transcriptions. Then, we transcribed all instances of query words within

Language	mHuBERT DTW			MPR VSM			MPR String Search		
	Top 5	MAP@5	MAP	Top 5	MAP@5	MAP	Top 5	MAP@5	MAP
Gujarati	44.0%	0.332	0.312	22.8%	0.163	0.148	19.9%	0.143	0.138
Hindi	43.6%	0.328	0.333	24.5%	0.167	0.177	27.0%	0.190	0.196
Marathi	61.8%	0.484	0.401	34.8%	0.265	0.201	28.1%	0.214	0.184
Odia	73.3%	0.517	0.432	30.0%	0.232	0.196	53.3%	0.367	0.304
Tamil	46.6%	0.357	0.340	12.7%	0.090	0.085	13.6%	0.089	0.086
Telugu	40.7%	0.307	0.286	18.9%	0.133	0.118	20.0%	0.142	0.131
Average	51.7%	0.388	0.351	24.0%	0.175	0.154	27.0%	0.191	0.173

Table 2: Baseline results for all MUCS languages.

Layer	Top 5	MAP@5	MAP
7	53.3%	0.282	0.237
8	57.2%	0.313	0.261
9	59.3%	0.316	0.265
10	50.4%	0.264	0.216
11	39.8%	0.207	0.163

Table 3: Average metrics for the mHuBERT model for all languages (MUCS and Gormati), for various layers.

Language	Top 5	MAP@5	MAP
Gormati	68.7%	0.496	0.228
Gujarati	48.1%	0.214	0.215
Hindi	50.3%	0.226	0.242
Marathi	64.0%	0.308	0.274
Odia	86.7%	0.367	0.324
Tamil	51.2%	0.401	0.377
Telugu	45.8%	0.199	0.196
Average (MUCS)	57.7%	0.286	0.271

Table 4: Results using the layer 9 of mHuBERT.

documents and quantified the mismatch between these and the “references” using phone mismatch rate (PMR).⁴ The results highlighted the MPR’s poor performance, revealing a PMR of 54%.

We also examined how well the MPR can detect voice activity. Using the 2 hours of Tamil training data and comparing it to the gold labels, we determined there were 89 minutes of voice activity. However, the MPR only detected 63 minutes.

5.2 Baseline DTW and MPR Results

We found that DTW consistently outperformed both types of retrieval using the MPR transcriptions, see Table 2. These results suggest that the MPR transcriptions were too inconsistent for our approximate string retrieval methods. Nevertheless, we use these results as baselines for comparison to our subsequent model results.

5.3 mHuBERT Results

The most important consideration for mHuBERT is what layer to extract the representations from. We hypothesised that representations from layers 8-10 would be most effective, due to results by Pasad et al. (2024). Average results per layer are shown in Table 3. We found that layer 9 was optimal across all languages. Table 4 shows results for all languages, for layer 9. For all following experiments, we use mHuBERT features from layer 9.

⁴Phone mismatch rate is similar to phone error rate.

5.4 AWE Model

After a series of initial experiments, we determined an optimal inference length of 3-9 phones for both time and phone (gold and MPR) window inference for the MUCS languages. We additionally found optimal values of: 9, 0.07 and 10^{-4} , for layer, temperature and learning rate, respectively.

Note that the optimal inference length of 3-9 phones varied slightly with training language and test language, although, 3-9 phones was either optimal or very near it. Also, for the MUCS languages, phone window (MPR) inference performed slightly better than time window inference. We use MPR inference with 3-9 phones on the MUCS languages for all following experiments. Results for these initial experiments are included in Appendix A.

We hypothesised the optimal inference length for Gormati to differ with that for the MUCS languages due to differences in the data and search task. For Gormati, our results indicated optimal lengths of 4-13 phones and 3-7 phones for time window and phone window (MPR), respectively. The time window length is much longer than that for the MUCS languages. This could be because Gormati queries are generally much longer (average 34 s) than our MUCS language queries (average <1 s), meaning longer phone sequences occur more frequently and therefore may be more discriminative. In contrast, the phone window (MPR) length is similar to that

Language	Inference	Top 5	MAP@5	MAP
Gormati	MPR	73.7%	0.566	0.252
	Time	88.9%	0.670	0.310
MUCS Average	Gold	75.1%	0.393	0.388
	MPR	67.2%	0.345	0.335
	Time	64.7%	0.338	0.329

Table 5: Results for Gormati for the Tamil-trained model and MUCS for the ensemble model with time and phone (Gold, MPR) window inference.

for the MUCS languages. This could be since the MPR is inaccurate, regularly deletes phones and inserts silences, meaning long phone sequences are less likely to occur and those that do occur are unlikely to be transcribed correctly.

Table 5 shows that on Gormati, time window inference performs much better than phone window (MPR) inference. We tested with the Tamil-trained model but results are similar for the other models. This is in contrast to the MUCS languages, where phone window (MPR) inference slightly outperforms time window inference. This could indicate that the MPR performs much worse on Gormati than other languages. For all subsequent experiments, Gormati inference is done with a time window with length 4-13 phones.

5.5 Ensemble Models

We hypothesized that by ensembling models trained on different languages we may produce a model that performs well over all metrics and over all languages. To ensemble models, we simply averaged the scores for each document, for each query, over all models. Results in Table 6 show that the ensemble model performs well over all metrics.

Considering inference methods again, Table 5 shows that for the ensemble model, phone window (MPR) inference performs on average slightly better than time window inference for the MUCS languages, matching our initial results. A full breakdown is in appendix B.

Table 5 shows that the performance of phone window (MPR) inference is much lower than the top-line results with the gold labels. This highlights the importance of a good MPR and demonstrates that performance can still be greatly enhanced by improving the MPR.

5.6 Training with MPR Labels

Training with MPR-predicted phone timings removes the requirement for labelled training data,

Language	Top 5	MAP@5	MAP
Gormati	87.9%	0.683	0.336
Gujarati	62.9%	0.286	0.291
Hindi	60.1%	0.268	0.284
Marathi	69.7%	0.357	0.333
Odia	90.0%	0.410	0.371
Tamil	61.5%	0.479	0.465
Telugu	59.0%	0.270	0.265
Average (MUCS)	67.2%	0.345	0.335

Table 6: Results using the ensemble AWE model. For the MUCS languages, MPR phone window inference is used and for Gormati, time window is used.

Language	Labels	Top 5	MAP@5	MAP
Gormati	Gold	87.9%	0.683	0.336
	MPR	82.8%	0.660	0.315
MUCS Average	Gold	67.2%	0.345	0.335
	MPR	62.2%	0.325	0.314

Table 7: Results with the ensemble model, training using MPR or Gold labels.

which is useful as low-resource languages often lack labelled data. We expect that training using MPR-predicted phone timings will produce a worse model than using gold timings, due to the added noise from the MPR. However, as the MPR is reasonably effective for inference, we expect that a model trained with MPR-timings could still be reasonably effective. To test this, we retrained our models using the MPR-predicted timings and compared it to our previous models trained on gold labels. Table 7 shows that the ensembled MPR-trained models are clearly worse than the gold label-trained models, as expected. However, they still perform reasonably well, with metrics that are only at most 7% lower than those of the gold labelled models. These results show that labels are not necessary for building a strong model, and a fully unsupervised transfer learning approach using an MPR can be effective.

5.7 Training with Gormati Data

We hypothesised that training with Gormati (using MPR-predicted labels) could improve performance as we train with the same language we test on. However, based on the results in Section 5.6, it seems like the MPR may not perform well on Gormati. To test this, we finetuned our gold label trained Tamil model on Gormati using files previously excluded from the search collection, totalling around 30 minutes of audio. We used Tamil since

it had the highest Top 5 score on Gormati.

We found that finetuning with Gormati degrades model performance. We could have potentially tested further by partitioning additional Gormati data from the search collection and finetuning the model’s hyperparameters. However, we chose not to do this since the initial results were very poor and indicated that this method would be unsuccessful.

5.8 Amount of Data

In low-resource contexts, it is useful to know how much data is necessary to train a model effectively. Lower data requirements could enable the use of data from languages that are more closely related to the target language, even if they have less data than other higher-resource but less related languages.

All previous models were trained using 2 hours of data. Here we tested the effect of training using half and quarter that amount. We tested using the Tamil model since it had the best Top 5 score on Gormati. We found that reducing the training data to 1 hour produces a very similar model to 2 hours. Further reducing the data to 0.5 hours noticeably impacts performance, though not too drastically. Therefore, in general, increasing the training data increases performance, with the greatest increase between 0.5 to 1 hour of data.

6 Discussion

The best results for each model on the MUCS languages are shown in Table 8. The AWE model has the best average Top 5 score with 67.2%, though it has a slightly worse MAP and MAP@5 score compared to the mHuBERT DTW baseline. This suggests that the AWE model is much better at producing at least one correct response per query than the mHuBERT DTW model but it is slightly worse when it comes to the overall ranking. Combining these two models could produce a model with high scores over all metrics.

On Gormati, the AWE model performs the best with a Top 5 score of 88.9%, exceeding the best score of 74% from [Reitmaier et al. \(2024\)](#). note that the mHuBERT DTW model cannot readily be applied to the more complex Gormati document retrieval task. Unlike [Reitmaier et al. \(2024\)](#), our model requires no target language training data and thus can operate on classes that have just one document. Similarly, the success of our transfer learning approach shows that this method can be extended to other low-resource languages using only

Model	Top 5	MAP@5	MAP
AWE Ensemble	67.2%	0.345	0.335
mHuBERT	57.7%	0.286	0.271
mHuBERT DTW	52%	0.39	0.35

Table 8: Average MUCS language results with a selection of the best-performing models.

one hour of data from a related higher-resource language. For best performance, this data must be labelled. However, we showed that a good model can be trained using MPR labels. This shows that a successful model can be produced without any supervised data from the target language, a critical requirement for an unwritten language.

The success of the transfer learning approach suggests that the AWEs are mainly capturing phonetic information, as semantic information is unlikely to generalise well across languages. A method that captures more semantic information might produce better results given the nature of the Gormati data. [Jacobs and Kamper \(2024\)](#) present such a method, but it requires knowledge of word boundaries, making it unsuitable for unwritten languages. In the future, adapting this or similar methods to unwritten languages could increase the semantic content of the AWEs, potentially improving information retrieval further.

In a case where there is no data from related languages or resources are unavailable for training, then the mHuBERT or mHuBERT DTW models could be used. They perform worse than the AWE model, but require no training and so can be applied directly to the target language.

7 Conclusion

We have presented a successful unsupervised method for developing a purely speech-based IR system. However, there remain several avenues for future work. Improving our inference method by experimenting with window lengths, strides and overlaps could be valuable. Optimising model architectures could also lead to improvements, as might combining the AWE model with the mHuBERT DTW model. Our model development with the MUCS data was geared towards the specific task of returning documents that directly contained a single-word query. This type of retrieval is insufficient when it is necessary to return semantically similar results to the query, or for multi-word queries that might benefit from partial matching.

8 Limitations

There are several limitations to this study, all of which can be addressed with further work. First, we only experimented with training using Tamil, Telugu and Gujarati, as these were the only related languages where we had approximately similar speech data. However, with additional data, it would be possible to train using other languages more closely related to Gormati, such as Marathi and Hindi. We did not experiment with multilingual training, which could enhance the models' ability to generalise to other languages; neither did we train mHuBERT on related languages to improve the quality of its representations. Our document ranking system was not tuned to the Gormati search task; in future, we could experiment with different similarity metrics and different methods to compare queries and documents. We only experimented with mHuBERT representations but we could experiment with a wider range of self-supervised representations to better determine the optimal representation. We used a somewhat limited number of documents and queries for testing on Odia, Hindi and Marathi; increasing the number of queries and documents would increase our confidence of our results with these languages. Finally, the Gormati dataset used in this work was designed with IR in mind, and was collected collaboratively with members of the Banjara community. When applying methods from this work to other languages, especially low-resource languages, it is important to keep in mind the community being served. This could take the form of catering the system towards a specific application or domain that is most useful to speakers of the target language, or involving speakers in the evaluation process.

Acknowledgments

References

Robin Algayres, Adel Nabli, Benoît Sagot, and Emmanuel Dupoux. 2022. [Speech Sequence Embeddings using Nearest Neighbors Contrastive Learning](#). In *Proc. Interspeech 2022*, pages 2123–2127.

Arnon Amir, Alon Efrat, and Savitha Srinivasan. 2001. Advances in phonetic word spotting. In *Proceedings of the tenth international conference on information and knowledge management*, pages 580–582.

Marcely Zanon Boito, Vivek Iyer, Nikolaos Lajos, Laurent Besacier, and Ioan Calapodescu. 2024. [mHuBERT-147: A Compact Multilingual HuBERT Model](#). *arXiv preprint*. ArXiv:2406.06371 [cs, eess].

Andi Buzo, Horia Cucu, Mihai Safta, and Corneliu Burileanu. 2013. [Multilingual query by example spoken term detection for under-resourced languages](#). In *2013 7th Conference on Speech Technology and Human - Computer Dialogue (SpED)*, pages 1–6, Cluj-Napoca, Romania. IEEE.

Ciprian Chelba, Timothy J. Hazen, and Murat Sarraciar. 2008. [Retrieval and browsing of spoken content](#). *IEEE Signal Processing Magazine*, 25(3):39–49. Conference Name: IEEE Signal Processing Magazine.

Anni R Coden, Eric W Brown, and Savitha Srinivasan. 2002. *Information retrieval techniques for speech applications*, volume 2273. Springer Science & Business Media.

Ana Deumert. 2014. [Sociolinguistics and Mobile Communication](#). Edinburgh University Press.

Anuj Diwan, Rakesh Vaideeswaran, Sanket Shah, Ankita Singh, Srinivasa Raghavan, Shreya Khare, Vinit Unni, Saurabh Vyas, Akash Rajpuria, Chiranjeevi Yarra, Ashish Mittal, Prasanta Kumar Ghosh, Preethi Jyothi, Kalika Bali, Vivek Seshadri, Sunayana Sitaram, Samarth Bharadwaj, Jai Nanavati, Raoul Nanavati, and Karthik Sankaranarayanan. 2021. [MUCS 2021: Multilingual and Code-Switching ASR Challenges for Low Resource Indian Languages](#). In *Interspeech 2021*, pages 2446–2450. ISCA.

Toni Giorgino. 2009. [Computing and Visualizing Dynamic Time Warping Alignments in R: The dtw Package](#). *Journal of Statistical Software*, 31(7):1–24.

Trina Gorman, Emma Rose, Judith Yaaqoubi, Andrew Baylor, and Beth Kolko. 2011. [Adapting usability testing for oral, rural users](#). In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, pages 1437–1440, Vancouver BC Canada. ACM.

Kira Griffitt and Stephanie Strassel. 2016. The query of everything: developing open-domain, natural-language queries for bolt information retrieval. In *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC'16)*, pages 3741–3747.

William Hartmann, Damianos Karakos, Roger Hsiao, Le Zhang, Tanel Alumäe, Stavros Tsakalidis, and Richard Schwartz. 2017. Analysis of keyword spotting performance across iarpa babel languages. In *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 5765–5769. IEEE.

Timothy J. Hazen, Wade Shen, and Christopher White. 2009. [Query-by-example spoken term detection using phonetic posteriorgram templates](#). In *2009 IEEE Workshop on Automatic Speech Recognition & Understanding*, pages 421–426.

732	Wei-Ning Hsu, Benjamin Bolte, Yao-Hung Hubert Tsai,	Thomas Reitmaier, Dani Kalarikalayil Raju, Ondrej Kle-	788
733	Kushal Lakhotia, Ruslan Salakhutdinov, and Abdel-	jch, Electra Wallington, Nina Markl, Jennifer Pear-	789
734	rahman Mohamed. 2021. HuBERT: Self-Supervised	son, Matt Jones, Peter Bell, and Simon Robinson.	790
735	Speech Representation Learning by Masked Predic-	2024. Cultivating Spoken Language Technologies	791
736	tion of Hidden Units . <i>IEEE/ACM Transactions on</i>	for Unwritten Languages . In <i>Proceedings of the CHI</i>	792
737	<i>Audio, Speech, and Language Processing</i> , 29:3451–	<i>Conference on Human Factors in Computing Systems</i> ,	793
738	3460.	pages 1–17, Honolulu HI USA. ACM.	794
739	Yushi Hu, Shane Settle, and Karen Livescu. 2021.	F. Richardson, M. Ostendorf, and J.R. Rohlicek. 1995.	795
740	Acoustic Span Embeddings for Multilingual Query-	Lattice-based search strategies for large vocabulary	796
741	by-Example Search . In <i>2021 IEEE Spoken Language</i>	speech recognition. In <i>1995 International Confer-</i>	797
742	<i>Technology Workshop (SLT)</i> , pages 935–942, Shen-	<i>ence on Acoustics, Speech, and Signal Processing</i> ,	798
743	zhen, China. IEEE.	volume 1, pages 576–579 vol.1.	799
744	Christiaan Jacobs and Herman Kamper. 2021. Mul-	Nay San, Martijn Bartelds, Mitchell Browne, Lily Clif-	800
745	tilingual Transfer of Acoustic Word Embeddings	ford, Fiona Gibson, John Mansfield, David Nash,	801
746	Improves When Training on Languages Related to	Jane Simpson, Myfany Turpin, Maria Vollmer, Sasha	802
747	the Target Zero-Resource Language . In <i>Interspeech</i>	Wilmoth, and Dan Jurafsky. 2021. Leveraging Pre-	803
748	2021, pages 1549–1553. ISCA.	Trained Representations to Improve Access to Un-	804
749	Christiaan Jacobs and Herman Kamper. 2024. Leverag-	transcribed Speech from Endangered Languages . In	805
750	ing multilingual transfer for unsupervised semantic	<i>2021 IEEE Automatic Speech Recognition and Un-</i>	806
751	acoustic word embeddings . <i>IEEE Signal Processing</i>	<i>derstanding Workshop (ASRU)</i> , pages 1094–1101.	807
752	<i>Letters</i> , 31:311–315.	Ramon Sanabria, Ondřej Klejch, Hao Tang, and Sharon	808
753	David A James and Steve J Young. 1994. A fast lattice-	Goldwater. 2023a. Acoustic Word Embeddings for	809
754	based approach to vocabulary independent wordspot-	Untranscribed Target Languages with Continued Pre-	810
755	ting. In <i>Proceedings of ICASSP'94. IEEE Interna-</i>	training and Learned Pooling . In <i>INTERSPEECH</i>	811
756	<i>tional Conference on Acoustics, Speech and Signal</i>	2023, pages 406–410. ISCA.	812
757	<i>Processing</i> , volume 1, pages I–377. IEEE.	Ramon Sanabria, Hao Tang, and Sharon Goldwa-	813
758	Xinjian Li, Siddharth Dalmia, Juncheng Li, Matthew	ter. 2023b. Analyzing acoustic word embeddings	814
759	Lee, Patrick Littell, Jiali Yao, Antonios Anastasopou-	from pre-trained self-supervised speech models . In	815
760	los, David R. Mortensen, Graham Neubig, Alan W	<i>ICASSP 2023 - 2023 IEEE International Confer-</i>	816
761	Black, and Florian Metze. 2020. Universal Phone	<i>ence on Acoustics, Speech and Signal Processing</i>	817
762	Recognition with a Multilingual Allophone System .	(ICASSP), pages 1–5.	818
763	In <i>ICASSP 2020 - 2020 IEEE International Con-</i>	P Sudhakar, K Sreenivasa Rao, and Pabitra Mitra. 2023.	819
764	<i>ference on Acoustics, Speech and Signal Process-</i>	Query-by-Example Spoken Term Detection for Zero-	820
765	<i>ing (ICASSP)</i> , pages 8249–8253, Barcelona, Spain.	Resource Languages Using Heuristic Search . <i>ACM</i>	821
766	IEEE.	<i>Transactions on Asian and Low-Resource Language</i>	822
767	Nicolas Moreau, Hyoung-Gook Kim, and Thomas	<i>Information Processing</i> .	823
768	Sikora. 2004. Combination of phone n-grams for	Paolo Tormene, Toni Giorgino, Silvana Quaglini, and	824
769	a mpeg-7-based spoken document retrieval system.	Mario Stefanelli. 2009. Matching incomplete time	825
770	In <i>2004 12th European Signal Processing Confer-</i>	series with dynamic time warping: an algorithm and	826
771	<i>ence</i> , pages 549–552.	an application to post-stroke rehabilitation . <i>Artificial</i>	827
772	Kenney Ng and Victor W. Zue. 2000. Subword-based	<i>Intelligence in Medicine</i> , 45(1):11–34.	828
773	approaches for spoken document retrieval . <i>Speech</i>	Ewald van der Westhuizen, Herman Kamper, Raghav	829
774	<i>Communication</i> , 32(3):157–186.	Menon, John Quinn, and Thomas Niesler. 2022. Fea-	830
775	Joseph Olive, Caitlin Christianson, and John McCary.	ture learning for efficient ASR-free keyword spotting	831
776	2011. <i>Handbook of natural language processing</i>	in low-resource languages . <i>Computer Speech and</i>	832
777	<i>and machine translation: DARPA global autonomous</i>	<i>Language</i> , 71(C).	833
778	<i>language exploitation</i> . Springer Science & Business	Michael Weintraub. 1993. Keyword-spotting using sri's	834
779	Media.	decipher large-vocabulary speech-recognition system.	835
780	Alex Park and James Glass. 2008. Unsupervised Pattern	In <i>1993 IEEE International Conference on Acoustics,</i>	836
781	Discovery in Speech . <i>IEEE Transactions on Audio,</i>	<i>Speech, and Signal Processing</i> , volume 2, pages 463–	837
782	<i>Speech, and Language Processing</i> , 16:186–197.	466. IEEE.	838
783	Ankita Pasad, Chung-Ming Chien, Shane Settle, and	Yaodong Zhang et al. 2013. <i>Unsupervised speech pro-</i>	839
784	Karen Livescu. 2024. What Do Self-Supervised	<i>cessing with applications to query-by-example spo-</i>	840
785	Speech Models Know About Words? <i>Transactions</i>	<i>ken term detection</i> . Ph.D. thesis, Massachusetts In-	841
786	<i>of the Association for Computational Linguistics</i> ,	stitute of Technology.	842
787	12:372–391.		

A Single System AWE Results

We trained our AWE models on three languages: Tamil, Telugu, and Gujarati. Table 9 shows results for each test languages using MPR phone inference for MUCS languages and time window inference for Gormati.

B AWE Ensemble Model Results

Table 10 contains a breakdown of the results for the ensemble model, trained on gold labels over different inference methods and MUCS test languages.

Language	Tamil training			Telugu training			Gujarati training		
	Top 5	MAP@5	MAP	Top 5	MAP@5	MAP	Top 5	MAP@5	MAP
Gormati	88.9%	0.670	0.310	85.9%	0.696	0.336	85.9%	0.686	0.343
Gujarati	57.7%	0.261	0.265	58.4%	0.261	0.268	63.5%	0.286	0.289
Hindi	55.2%	0.244	0.260	56.4%	0.251	0.267	62.0%	0.274	0.286
Marathi	65.2%	0.334	0.314	65.2%	0.333	0.314	66.3%	0.340	0.307
Odia	80.0%	0.383	0.350	80.0%	0.376	0.339	73.3%	0.368	0.350
Tamil	59.7%	0.456	0.442	59.4%	0.451	0.436	59.8%	0.454	0.439
Telugu	54.3%	0.251	0.246	57.3%	0.263	0.259	55.8%	0.257	0.252
<i>Average (MUCS)</i>	<i>62.0%</i>	<i>0.332</i>	<i>0.313</i>	<i>62.8%</i>	<i>0.323</i>	<i>0.314</i>	<i>63.5%</i>	<i>0.330</i>	<i>0.321</i>

Table 9: Results on each test language (using MPR phone inference for MUCS languages and time window inference for Gormati) for AWE models with different training languages. The average is only shown over MUCS languages.

Language	Phone Window (Gold)			Phone Window (MPR)			Time Window		
	Top 5	MAP@5	MAP	Top 5	MAP@5	MAP	Top 5	MAP@5	MAP
Gujarati	69.8%	0.315	0.319	62.9%	0.286	0.291	60.6%	0.274	0.277
Hindi	70.6%	0.315	0.332	60.1%	0.268	0.284	59.5%	0.277	0.289
Marathi	79.8%	0.397	0.377	69.7%	0.357	0.333	70.8%	0.352	0.337
Odia	93.3%	0.473	0.450	90.0%	0.410	0.371	76.7%	0.388	0.362
Tamil	71.1%	0.562	0.553	61.5%	0.479	0.465	61.2%	0.469	0.452
Telugu	66.1%	0.293	0.296	59.0%	0.270	0.265	59.2%	0.266	0.259
<i>Average</i>	<i>75.1%</i>	<i>0.393</i>	<i>0.388</i>	<i>67.2%</i>	<i>0.345</i>	<i>0.335</i>	<i>64.7%</i>	<i>0.338</i>	<i>0.329</i>

Table 10: Results for ensemble model for different inference methods.