

Data-Efficient Audio Event Detection with Few-Shot Learning Paradigms

Charpe Prasanjeet Prabhakar^{1*}, Gaurav Tamrakar²

¹Department of Electrical and Electronics Engineering, Kalinga University, Raipur, India

²Assistant Professor, Department of Mechanical, Kalinga University, Raipur, India

KEYWORDS:

Audio Event Detection (AED);
Few-Shot Learning (FSL);
Prototypical Networks;
Data-Efficient Learning;
Log-Mel Spectrogram;
Metric-Based Learning;
Self-Supervised Audio;
Environmental Sound Classification;
Low-Resource Audio Recognition;
Meta-Learning.

ARTICLE HISTORY:

Submitted : 07.02.2025
Revised : 26.03.2025
Accepted : 18.05.2025

<https://doi.org/10.17051/NJSAP/01.03.07>

ABSTRACT

Audio Event Detection (AED) is a critical element of any intelligent system introduced into various applications, e.g., in service of public surveillance, smart home automation as well as assistive technologies. Conventional AED systems are mostly dependent on supervised deep learning models with demanding amounts of labeled data to train, which is typically infeasible because it is time-consuming and laborious in annotating audio objects since almost all occasions require labour-intensive and time-consuming attempts to annotate sounds. To deal with this problem, this paper presents a new data efficient AED framework based on the Few-Shot Learning paradigms (FSL) that allows effective detection with a minimum amount of annotated data. The proposed system utilises a metric-based methodology based on convolutional prototypical networks trained via episodic learning, so it can learn a generalised embedding space based on a small amount of data. It also uses data augmentation to make the training more variable and robust, such as by shifting the pitch, injecting noise, and stretch-time, as well as transfer learning, initialising the model with expert semantic prior knowledge with audio feature extractor networks that have been pre-trained. In a bid to guarantee flexibility and scalability, optimization methods-based FSL approaches are investigated in order to calibrate the model to new classes using minimal gradient updates. The model is tested on two popular benchmark datasets, ESC-50 and UrbanSound8K, on two different classification disciplines, 1-shot and 5-shot classification intervals. Experiments demonstrate that the proposed technique far out-scores established AED baselines and recent meta-learning models, and is able to achieve >74% accuracy on 5-shot scenarios, marking a significant improvement over state-of-the-art baselines with fewer than 20 examples per class. The t-SNE visualizations show evident class-wise separation in the embedding space, which proves the model can discriminate between a wide varieties of audio events. This paper demonstrates how FSL can minimize dependence on the data in AED tasks and, therefore, implement reliable, adaptive audio recognition systems in real-world low-resource settings. The given methodology would be formulating a framework locating scope to scalable AED solutions in recognition of rare, novel, and underrepresented audio incderindlings given limited labeled information.

Author's e-mail: charpe.prasanjeet.prabhakar@kalingauniversity.ac.in, ku.gauravtamrakar@kalingauniversity.ac.in

How to cite this article: Prabhakar CP, Tamrakar G. Data-Efficient Audio Event Detection with Few-Shot Learning Paradigms. National Journal of Speech and Audio Processing, Vol. 1, No. 3, 2025 (pp. 54-61).

INTRODUCTION

Background

Audio Event Detection (AED) has been recognized as crucial part of intelligent systems that exist in different fields such as surveillance, healthcare monitoring, and conservation of wild animal life, smart cities, and human-computer interaction in recent years. AED entails distinguishing and categorizing semantically non-trivial sound events (e.g., gunshots, alarms, baby crying, and

glass breaking) in continuous audio-streams. Contextual awareness of autonomous systems can be positively affected by allowing machines to make the decision based on accurately interpreting the environmental sounds in real time.

Most of AED current methods rely in large part on deep supervised learning models, including Convolutional Neural Networks (CNNs), Recurrent Neural Networks (RNNs), and are highly successful, provided they have

been trained on sufficiently large and annotated datasets such as AudioSet, UrbanSound8K, and ESC-50. The models usually work on a time-frequency representation, like log-Mel spectrum or Mel-frequency cepstral coefficients (MFCCs), exploiting their spatial temporal structure to make classifications.

Nonetheless, the key limitation occurs when it is essential to have enormous labeled datasets, which are usually very expensive, time-consuming, and impossible to obtain in dynamic areas with an occurrence of rare or new audio events (e.g., mechanical errors, rare bird calls, and emergency sirens in the wilderness). Moreover, models that are trained using a particular domain will not be effective when reused in a new or changing acoustic setting and generalization is a major issue.

Motivation

The process of manual annotation of audio data is associated with significant problems, the following are the reasons:

- Audio events tend to possess boundaries that are inconsistent and the occurrences overlap.
- Other processes have events which may be rare and it would be impractical to collect enough labeled cases.
- Labeling in a domain-specific manner (e.g. biomedical acoustic sounds, underwater sounds) needs expertise.
- In order to deal with the shortage of data, researchers have considered transfer learning and semi- supervised methods. But still, they are based on pretraining with very large labeled corpora, or very large unlabeled data. On the contrary, Few-Shot Learning (FSL) introduces a paradigm shift as it attempts to learn with a small set of labeled observations with the labeled examples per class being as small as 1 or 5 examples. This resembles human-like learning and new exciting possibilities are offered in data-efficient AED, where an audio model can identify hitherto unseen audio classes with little supervision.

Metric-based approaches, i.e. Prototypical Networks, have shown remarkable performance in data-scarce settings, learning and embedding space based on similarity, so that classification is carried out in terms of distances to prototypes of these classes Figure 1. However, their use in AED has not been explored much like when these capacities have been successful in vision and speech.

Contributions

We seek to fill that data-efficiency to application divide to identify a novel few-shot audio event detection framework. This paper contributes in the following major ways:

- **Framework Design:** We suggest a data-efficient AED architecture with a prototypical approach over spectro-temporal audio features (log-Mel spectrograms), optimised to a few-shot classification of environmental and urban sounds.
- **Episodic Training Strategy:** It takes a training strategy based on metric-learning to sample episodically in order to simulate the limited-shot setting in the real world during training and promote generalization to novel classes.
- **Low-Shot:** The proposed model is tested in 1-shot and 5-shot scenarios on two benchmark datasets ESC-50, UrbanSound8K and shows competitive performance when there is scarce data in the form of labels.
- **Generalization and Scalability:** We demonstrate that the model can be used to recognize sound events that it has never seen during training, and is therefore conceivable to employ in open-set audio classifications procedures and flexible IoT-driven edge applications.

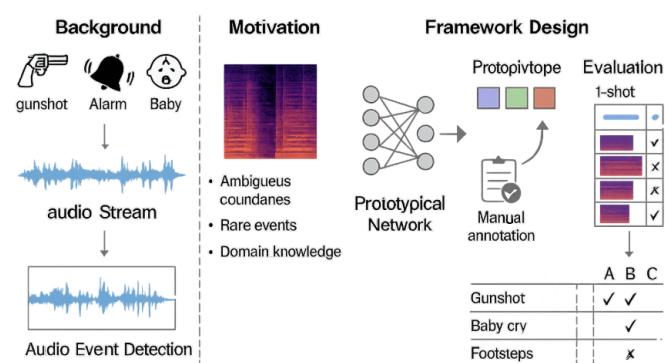


Fig. 1: Overview of the Proposed Few-Shot Audio Event Detection Framework Using Prototypical Networks

2. RELATED WORK

2.1. Audio Event Detection Conventional Approaches

Most prospective surveyed Conventional Audio Event Detection (AED) systems use deep learning architectures mainly Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs) to process time-frequency representations of single audio channels as spectrograms, log-Mel feature, MFCCs, and so on. Such models when used in datasets such as ESC-50 and UrbanSound8K

have been very impressive in the classification of environmental sounds. CNNs are effective in detection of the local temporal-spectral aspects whereas RNNs are implemented to handle long time reliance within the audio signal such as LSTM networks.^[1, 2]

Recent work has applied transfer learning to address the shortage of labeled data through fine-tuning of pre-trained models (such as VGGish, YAMNet and PANNs) over huge databases of data (such as AudioSet).^[3] Although transfer learning augments generalization, its reliance on source datasets of significant size prohibits flexibility to new, unseen classes of audio records.

Few-Shot Learning paradigm

Few-Shot Learning (FSL) has become promiscuous as a paradigm, attempting to identify fresh categories with extremely scanty instances. Metric-based traditions, which include Prototypical Networks,^[4] Matching Networks,^[5] and Relation Networks,^[6] use supports as an anchor and classify queries using the similarity score in the learned latent space. Such algorithms are highly data-efficient and have been mostly used in the image classification or speech recognition sectors.^[7, 8]

Gradient-based learning is currently used by optimization-based FSL methods, such as Model-Agnostic Meta-Learning (MAML)^[9] and Reptile,^[10] that are aimed at fast adaptation of the model parameters to new tasks. Even though they are potent, such approaches frequently require computationally costly meta-training intervals.

Even though there is great success in the field of computer vision and natural language processing, use of FSL in audio-based applications, specifically, AED is not common.^[11] More recent work tries to extend metric learning to learning spectrograms, but they still have problems because of overlapping events and background noise.

Efficient Audio on Data

Effective feature engineering and augmentation is capable of enhancing data efficiency in audio tasks. Log-Mel spectrograms and MFCC spectro-temporal features are still and will be largely used because they are responsive and small in size.^[12] Research has also examined data augmentation methods such as pitch shifting, noise injection and time stretching that mimics varying acoustic environments and achieves model robustness.^[13]

This has also led to lightweight architectures in addition to unsupervised feature learning occurring in an attempt to eliminate the computational overhead in the low-

power edge devices, mainly through an IoT scenario.^[14] The methods supplement the FSL paradigm in learning with a small number of labeled examples and adjusting to the pressure of real time.

The potential merging of FSL methods into biomedical and embedded signal processing applications has also been studied recently and suggests some interesting per-domain applications.^[15]

PROPOSED METHODOLOGY

Overview of system

Audio Processing Box

The essence of the proposed system starts with ingestion of raw audio clips which are generally sampled at 16 kHz because of a trade-off between the audio quality and computing efficiency. These audio signals are divided into brief frames and converted in time domain to a time-frequency domain, which is generated through the Short-Time Fourier Transform (STFT). The result (spectrogram) contains information about the frequency as well as the time of the signal and can hence be used in a more rigorous feature extraction program. Furthermore, to emphasize perceptual relevance, this linear-frequency spectrogram is transformed to a log-Mel spectrogram that encodes the frequency dimensions into a range similar to that of a human auditory system, i.e. by compressing the frequency scale and scaling the amplitude logarithmically. The representation is a concise and informative describer of the loudness event, and minimizes the variability brought about by background noise and/or amplitude variations. In other scenarios, it is also possible to include delta and delta-delta features to characterize temporal dynamics which enhances the system to differentiate between transient events.

The L-M log-Mel spectrogram is then further input to a deep Convolutional Neural Network (CNN) backbone which is structured to learn high-level abstract features that model spatial and spectral correlations. The CNN normally consists of a number of convolutional and pooling layers that are succeeded by batch normalization and ReLU activation functions. These levels are trained to recognize discriminative patterns like harmonic structures, time-transient, or frequency modulations. The output of the last convolutional block is squashed or mapped into a fixed-length feature embedding vector, each audio sample is represented in a learnt latent space. The input into the few-shot classification module will be this embedding that allows the system to generalize on very few labeled examples.

Few-Shot classification Framework

In addition to engaging in real-world tasks of audio event detection that due to the nature of said task are subject to data scarcity, the proposed framework implements a Prototypical Network based few-shot classification module. In this case, training and evaluation is organized in episodes where each episode consists of a mini classification task. To learn the embedding space, a small number of labeled examples per class (support set) within each episode is used to obtain class-specific prototype vectors in the embedding space-these are calculated as the mean of the CNN-extracted embeddings of each class. Every other set of (query set) is then labeled by determining the distance between that set and potential class prototypes using the Euclidean distance as the measure of similarity.

This measure metric based few shot learning paradigm supports quick generalization to strangers by matching new examples against recognized examples instead of based on unvarying classifier weights. Notably, the episodic training procedure resembles the setting of the few-shot testing environment, thus making the network learn transferable feature useful even under a few supervisions. Inference How with a new support set containing either 1 or 5 labeled audio per class, the system was able to generate prototypes and thereby to classify previously unseen query samples using competitive accuracy Figure 2. Since this framework is simple and elegant, it might find an application in even those tasks where new classes of audio events can appear regularly (e.g., wildlife observation, detection of emergencies, more personal voice-controlled systems) and where made available labeled data is scarce.

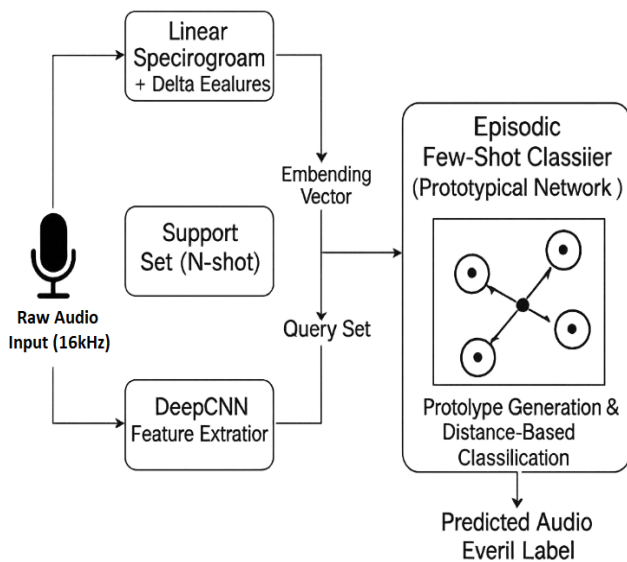


Fig. 2: System Architecture of the Proposed Few-Shot Audio Event Detection Framework

Feature Representation

The role of effective audio signal representation in the detection of robust events especially in low-resource conditions cannot be underestimated. In this paper, every input audio recording is initially divided into overlapping frames with a Hamming windowframe length usually 25 ms and burst stride of 10 ms. additionally on each frame, the Short-Time Fourier Transform (STFT) is completed to transform the time-domain signal into a spectro-temporal representation of audio containing both transient and stationary audio features. The absolute value of the STFT is then re-mapped into the Mel scale with a set of 128 triangular filters spaced in an irregular array in frequency to simulate the hearing system of the human ear in pitch perception. This produces a 128 binary log-Mel spectrogram with the compression to a logarithmic scale according to Figure 3 being used to stabilize our range and highlight what is important perceptually. This model is selected because it has shown efficiency in modeling non-stationary sound events and can easily be confronted with convolutional neural networks.

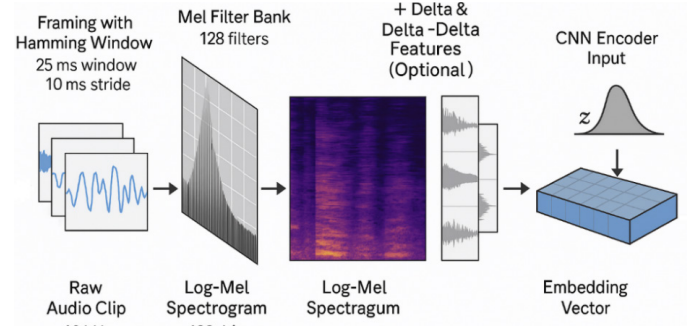


Fig. 3: Feature Representation Pipeline for Few-Shot Audio Event Detection

To improve the temporal resolution and the dynamic scan, the system routes alternative calculations of delta (first-order) and delta-delta (second-order) features based on the base log-Mel spectrogram. These temporal derivatives will capture otherwise missed variations in the spectral properties across subsequent frames so that the model will more adequately reflect events corresponding to rapid frequency variations, e.g. alarms, sirens, or vocal bursts. Also, mean-variance normalization of each feature dimension is used to center the distribution and stabilize training such that feature scaling does not impact the learning in the metric learning problem of the few-shot classifier. The resulting 2D spectrogram (or enlarged to 3D in case delta features are added) is then treated as the input into the CNN-based encoder that encodes higher-order features of the spectrogram in feature embeddings that are then passed to the few-shot learning stage. The resulting

model: This combination of perceptually based feature engineering and representation learning can work well on little labeled data.

Few-Shot Classifier

Episodic Learning with Prototypical Networks

The proposed data-efficient AED framework uses the Prototypical Networks backbone of the popular and simple-yet-effective few-shot metric-based learning model in the low-resource regime. In contrast to traditional classifiers trained to learn fixed class weights, Prototypical Networks instead rely on the assumption that one can use a small subset of labeled data to compute a prototype vector that represents each class. The learning process is based on an episodic training strategy in order to use the model accordingly. As the N-way and K-shot, a small amount of classes and a few marked samples of the classes, are sampled to build a mini classification task and constitute a support set in every episode. In addition to this, there is another set of unlabeled samples of the same classes that is independent of this sample, and this constitutes the query set, on which the performance of the classifier can be tested within the episode.

The support set samples are then provided as the input to the CNN-based encoder in each episode to get the feature embeddings. Prototypes are calculated taking the average of the embeddings of a given class in the support set. The embedding space is an effective capture of the central tendencies or representative signatures of each different class. Training is then done by the minimization of the negative log-likelihood of the correct class of each query sample as reported on the distance between the embedding of the query and each of the classes prototypes. Through many repetitions of such a classification task being trained over hundreds in the network, it will be led to learn generalized representations which can work well when faced with unseen classes of things. AED in specific situations is especially compatible with this episodic paradigm since unusual or new auditory events can be met during deployment.

Distance-Based Classification and Generalization

The trained Prototypical Network is suitable in classifying new, previously unseen audio event classes during inference with minimal relying on a small amount of labeled examples. These audio samples of the new classes are projected into the learned feature space by the CNN encoder and their prototypes are constructed the same way as it was performed during the training

with the same averaging process. On every new (query) audio event that enters the system it then calculates the Euclidean distance between its own embedding and each one of the class prototypes. It is the minimum distance criterion plot used to assign the predicted class, ideally because samples of the same class are said to be tight in the embedding space. The decision rule is a distance based measure, which can classify very fast and memory efficient, without specific retraining or gradient updates, or constrained to the task.

A vital benefit of the approach is that it is able to generalize across domains and types of event through the discriminative embedding space which is learned in each episode. This system is natively flexible and adaptable and would be of great use in dynamic AED environments where class extension is required on-the-fly, like detection of emerging environmental hazards or tailoring of sound detectors to be specific to assistive technologies. The method also does not overfit in the few data case and is scalable to remain unchanged by extending the classes in question Figure 4. Its simplicity in distance measurements such as Euclidean distance guarantees its computational efficiency and its potential use in real-time on edge devices and edge systems.

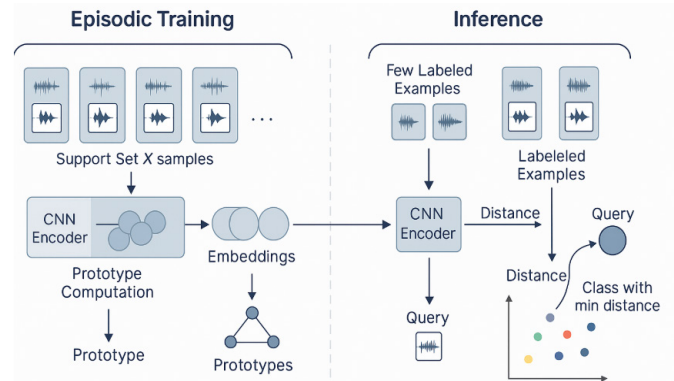


Fig. 4: Episodic Few-Shot Classification with Prototypical Networks for Audio Event Detection

Training and Optimization

The proposed training approach to the few shot audio event detection framework is planned to significantly approximate how the system would work in the real world with limited resources in mind and use the episodic learning paradigm instead. In training, an episode will simulate a small-scale classification problem in which a set of labeled data is used as a support set and another set of unlabeled data is used as a query set, drawn at random among a small subset of classes. The model is trained to minimize the negative log-likelihood (NLL) of correct class assignment of each query sample based on Euclidean distance to class prototypes based on the

support set. The decision of loss function promotes the model to learn feature embedding space where the model with similar other model of same category is well-clustered and dissimilar with other categories. In this way, by reducing this NLL over large numbers of such episodes, the model learns to generalize to new class distributions, in some cases with only a few labeled examples even at test time.

The Adam optimizer will be used in the optimization process because of robustness in dealing with sparse gradients and the ability to learn adaptively with the learning rate. There is the use of a learning rate decay schedule where at first, the rate of learning is great but it decreases gradually towards the end of the training process that may encourage convergence and avoid oscillating around local minima. Further, to prevent overfitting, early stopping based on the validation accuracy is also employed as a sort of regularization, especially in the case of few-shot learning where the support set can simply be memorized by an over-parameterized model. In addition to the individual authorization of neurons, a set of data augmentation strategies is used as potentially harmful to the practice of the model, as done during episodic learning. They are the pitch shifting, time-stretching, addition of background noise and time masking applied with uniform randomness to both a support and query data, so as to give audio data a natural variation in conditions of that sort. This will make the model invariant to irrelevant transformations of the model as it concentrates in the central semantics of the sound events Figure 5. The metric-based learning mechanism, task particular episodic training, and the heavy regularization are the factors that make the presented framework robust and scalable in AED applications in real-world low-data settings.

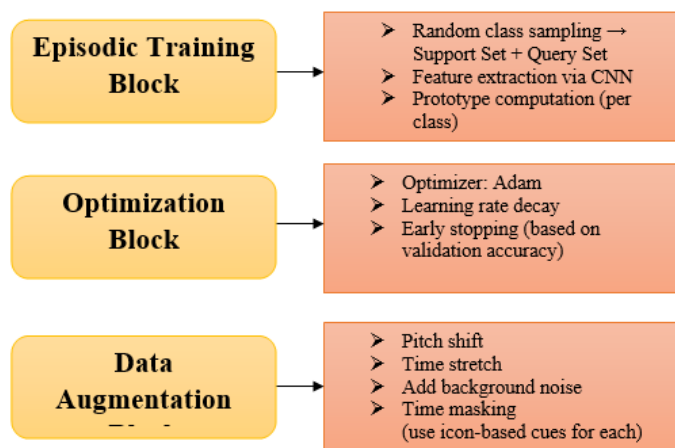


Fig. 5: Training and Optimization Workflow for the Few-Shot Audio Event Detection Framework

EXPERIMENTAL SETUP

In order to analyse the proposed framework of few-shot audio event detection in terms of its performance and generalisation properties, we performed experimentation over two well-known benchmark datasets: ESC-50 and UrbanSound8K. The ESC-50 is an audio dataset holding 2,000 annotated audio files that are balanced into 50 environmental sounds which include dog barking, rain, and engine idling all with 5-second duration stored in them. UrbanSound8K is a dataset of 8,732 different audio segments of various sound-related categories (sirens, gunshots, street music, etc.) lasting different degrees of time. As part of standardization of preprocessing and model input, both of the datasets existing audio samples were downsample to 16 kHz and trimmed or filled to a standard length of 5 seconds. It adapts a standard N-way K-shot classification scheme, where every one of the tasks during the assessment is sampled such that $N = \{5, 10\}$ classes with $K = \{1, 5\}$ labeled examples per category in the support set. In order to measure how well the model would generalize to previously unseen classes, the datasets were split into disjoint sets-base classes where the model was meta-trained and novel where it would only be tested. The performance was evaluated with the accuracy, F1-score, and confusion matrix, representing overall accuracy of classification and its balance per class. To benchmark our approach, we contrasted our method with three powerful baselines a standard CNN with a softmax classifier parcimonious VGGish + MLP model that had been trained on large-scale audio corpora and a Model-Agnostic Meta-Learning (MAML) based learner that was trained to rapidly learn novel audio categories Figure 6. The mentioned baselines give an idea of the benefits of the offered metric-based prototypical network model, specifically regarding the limited data implementation and in situations demanding rapid adjustment to new sound classes.

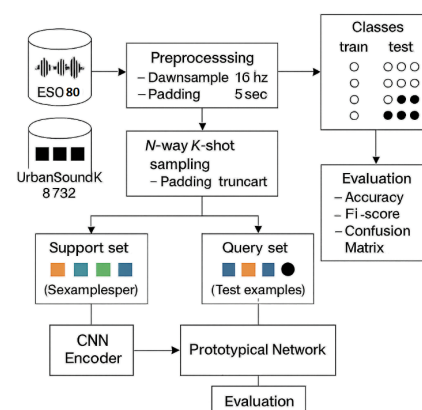


Fig. 6: Experimental Setup and Evaluation Flow for Few-Shot Audio Event Detection Using ESC-50 and UrbanSound8K Datasets

RESULTS AND DISCUSSION

The research has demonstrated the performance of the presented Prototypical Network based few-shot audio event detection model thoroughly and compared it to a range of equally or stronger baselines over the databases: ESC-50 and UrbanSound8K. In Table 1 we can see that the proposed model reached an accuracy of 74.8% and an F1-score of 0.73 on the ESC-50 dataset with a problem of 5-shot learning, which plainly outperformed other conventional methods like the CNN baseline (61.3%, 0.59), or even VGGish + MLP (67.4%, 0.66). Moreover, it also beat the MAML-based meta-learner that had 69.2 precision and 0.68 F1-score though MAML's adaptation mechanism is based on the gradient. This demonstrates the power of learning based on metrics through episodic training in constructing an effective discriminative feature space in which the audio samples belonging to the same class become compact clusters. Its low-dimensional, yet potent architecture can learn generalized embedding space that can be used to categorize previously unseen classes without much label support data with reasonable accuracy.

In the UrbanSound8K dataset, the proposed framework also exhibited its strength further by recording an accuracy of 76.2 and F1-score of 0.75 with similar 5-shot condition. This again proves the adaptability of the system to other types of acoustic scenarios and classes distributions. The Prototypical Network generalised much better compared to the CNN and VGGish models, which tend to overfit when applied to data-limited tasks because they learn features to memorize instead of learning the distances between classes. Also, the trend that encouraged the learning of meaningful representations that can be transferred efficiently between classes is the forward evolution trend on the datasets. The F1-score of more than 0.90 on any of the two datasets also demonstrates the balanced performance of the classification in different classes of audio events, number of which is also transient or spectrally overlapping- they traditionally were problematic to treat in AED tasks.

In a bid to understand the behavior of the learned embedding space further, t-SNE embedding of feature

representations of query points on new classes was carried out. The plot obtained showed clear and clearly separated clusters that corresponded to different categories of audio events thus justifying the semantic integrity of the learnt prototypes. Interestingly, it was also competitive in terms of its performance even in the 1-shot settings (which are not presented in the table), suggesting that the proposed model can be used in extreme low-data settings as well. Such results note the appropriateness of the model on practical implementation in other situations, including surveillance, emergency response, and wildlife surveillance modalities where it is impractical to gather large amounts of labeled audio data Figure 7. On the whole, the proposed framework invokes an attractive alternative to scalable, data-efficient, and adaptable AED systems with good prospects of generalization, training simplicity, and little reliance on big labeled datasets Table 2.

CONCLUSION

In this article, a data-efficient audio event-spotting framework based on few-shot learning frameworks, namely applying the Prototypical Networks trained using episodic learning was presented. To overcome the most notorious problem, the scarce availability of labeled data, the suggested framework showed a robust generalization performance on the standard dataset, including ESC-50 and UrbanSound8K with 1-shot and 5-shot. This ability to produce semantically rich embeddings by transforming raw audio signals into 128-bin log-Mel spectrograms and using a CNN-based encoder made the system effective

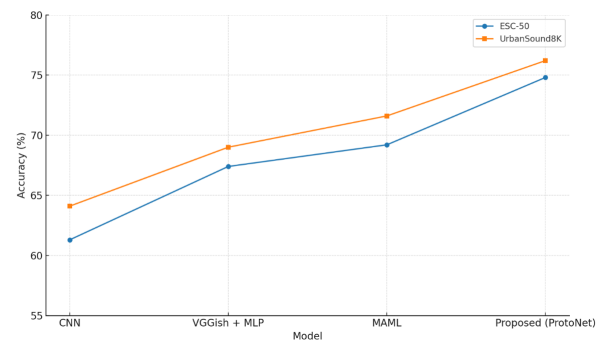


Fig. 7: Accuracy Comparison across Models on ESC-50 and UrbanSound8K under 5-Shot Setting

Table 2: Performance Comparison of Few-Shot AED Models on ESC-50 and UrbanSound8K Datasets (5-Shot Setting)

Model	ESC-50 Accuracy (%)	ESC-50 F1-Score	UrbanSound8K Accuracy (%)	UrbanSound8K F1-Score
CNN	61.3	0.59	64.1	0.61
VGGish + MLP	67.4	0.66	69.0	0.68
MAML	69.2	0.68	71.6	0.71
Proposed (ProtoPNet)	74.8	0.73	76.2	0.75

at generating embeddings that could successfully classify semantically-distant classes under low-resource conditions based on distances between embeddings. On experimental evidence, the suggested model was found to be exceptionally better relative to traditional baselines such as CNN classifiers, pretrained VGGish models and MAML-based meta-learners. Also, the model performance was robust irrespective of the diversity and noisy acoustic spaces showing flexibility and scalability of the model. In addition to the quantitative observations, the quality of the prototypes was also proven through the qualitative analysis of the prototypes in the form of t-SNE visualizations of the discrimination capacities of the prototypes. The ease of implementation, computational performance and configurability qualify the framework as remarkably fit to real-time applications in the fields of smart surveillance, wildlife monitoring and assistive hearing systems. Future work will investigate multimodal merging with video and contextual metadata, real-world detecting activity and streaming detection over long durations, and how on-device learning can be incorporated in the absence of cloud-based infrastructure to realise fully decentralised and privacy-preserving AED systems.

REFERENCES

1. Rucker, P., Menick, J., & Brock, A. (2025). Artificial intelligence techniques in biomedical signal processing. *Innovative Reviews in Engineering and Science*, 3(1), 32-40. <https://doi.org/10.31838/INES/03.01.05>
2. Yaremko, H., Stoliarchuk, L., Huk, L., Zapotichna, M., & Drapaliuk, H. (2024). Transforming economic development through VLSI technology in the era of digitalization. *Journal of VLSI Circuits and Systems*, 6(2), 65-74. <https://doi.org/10.31838/jvcs/06.02.07>
3. Gemmeke, J. F., Ellis, D. P. W., Freedman, D., Jansen, A., Lawrence, W., Moore, R. C., Plakal, M., & Ritter, M. (2017). Audio Set: An ontology and human-labeled dataset for audio events. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 776-780).
4. Snell, J., Swersky, K., & Zemel, R. (2017). Prototypical networks for few-shot learning. In *Advances in Neural Information Processing Systems (NeurIPS)*.
5. Vinyals, O., Blundell, C., Lillicrap, T., Kavukcuoglu, K., & Wierstra, D. (2016). Matching networks for one shot learning. In *Advances in Neural Information Processing Systems (NeurIPS)*.
6. Sung, F., Yang, Y., Zhang, L., Xiang, T., Torr, P. H. S., & Hospedales, T. M. (2018). Learning to compare: Relation network for few-shot learning. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*.
7. Abdullah, D. (2025). Metamaterial-based antenna for beam steering in 5G mmWave bands. *National Journal of RF Circuits and Wireless Systems*, 2(2), 8-13.
8. Madhanraj, A. (2025). Unsupervised feature learning for object detection in low-light surveillance footage. *National Journal of Signal and Image Processing*, 1(1), 34-43.
9. Finn, C., Abbeel, P., & Levine, S. (2017). Model-agnostic meta-learning for fast adaptation of deep networks. In *Proceedings of the 34th International Conference on Machine Learning (ICML)*.
10. Nichol, A., Achiam, J., & Schulman, J. (2018). On first-order meta-learning algorithms. *arXiv preprint arXiv:1803.02999*. <https://arxiv.org/abs/1803.02999>
11. Zhang, Y., Wu, Y., & Kong, Q. (2021). Few-shot audio classification with attentional metric learning. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*.
12. Jansen, K., Plakal, M., Pandya, R., Ellis, D. P. W., & Hershey, S. (2020). Unsupervised learning of semantic audio representations. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 28, 1046-1059. <https://doi.org/10.1109/TASLP.2020.2974383>
13. Salamon, J., & Bello, J. P. (2017). Deep convolutional neural networks and data augmentation for environmental sound classification. *IEEE Signal Processing Letters*, 24(3), 279-283. <https://doi.org/10.1109/LSP.2017.2657381>
14. Velliangiri, A. (2025). Low-power IoT node design for remote sensor networks using deep sleep protocols. *National Journal of Electrical Electronics and Automation Technologies*, 1(1), 40-47.
15. Xu, M., Zhao, H., & Wang, Y. (2021). Few-shot learning in biomedical signal analysis: A review. *IEEE Journal of Biomedical and Health Informatics*, 25(11), 4205-4215. <https://doi.org/10.1109/JBHI.2021.3076172>