



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 13 **Issue:** X **Month of publication:** October 2025

DOI: <https://doi.org/10.22214/ijraset.2025.74862>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

An Investigation of PLP Variants for Acoustic Scene Classification

Chandrasekhar Paseddula

Electronics and Communication Engineering G.Narayanamma Institute of Technology and Science, Hyderabad, Telangana, India

Abstract: In this paper, PLP coefficients and PLPCC features are investigated as a representation of an acoustic scene using DNN. We have experimented on DCASE 2018 Task 1 dataset and DCASE 2017 dataset. Experiments are carried out for subtasks A and B. We have experimented with individual feature sets as well as decision level DNN score fusion of different combinations of feature sets. From the experiments, it was observed that the proposed PLP and PLPCC give better performance for subtasks A and B. For subtasks A and B, individual PLP yields an improvement of 8.9% and 13.6% respectively. Further PLPCC resulted in an improvement of 8.6% and 12.5%. We have achieved significant improvements in accuracy for subtasks A (11.4%) and B (14.4%) after fusion of DNN decision level scores obtained from PLP, PLPCC and log mel-band energies compared to the 2018 baseline system. We have also experimented on 2017 dataset on 4 fold cross-validation, with individual PLP yielding an improvement of 5.8% and PLPCC achieving an improvement of 4.7%. The fusion of DNN decision level scores obtained from PLP, PLPCC and log mel-band energies gave an improvement of 6.0% compared to the 2017 baseline system.

Index Terms: Log-Mel band energies, Perceptual Linear Prediction (PLP), Acoustic Scene Classification (ASC), Deep Neural Network (DNN).

I. INTRODUCTION

The Acoustic Scene Classification (ASC) research in signal processing, machine learning and interdisciplinary fields has become more popular due to significance of information gathered from environmental sounds in various applications like surveillance, smartphones, robotics, data archiving, audio hearing aids, etc [2,3]. Initially, a large number of spectral, cepstral, energy and voicing-related audio features and SVM are used to classify these short segments and a majority voting scheme is employed in [4]. Spectral, temporal and spatial features with SVM classifier is used for ASC in [5]. Histogram of gradients of time-frequency representations for audio scene detection is investigated in [6]. A Bag-of-Features approach is used for acoustic event detection in [7]. Spectrogram pattern matching based environmental sound classification is used in [8]. Deep convolutional neural networks and data augmentation for environmental sound classification is carried out in [9]. Sound scene identification based on MFCC, binaural features and a support vector machine classifier is used for ASC in [10]. Frame based classification using Hidden Markov model was proposed in [11]. A hybrid approach using binaural I-vectors and deep convolutional neural networks is used for ASC in [12]. ASC using a CNN-super Vector system trained with auditory and spectrogram image features is used in [13]. Generative adversarial network based acoustic scene training set augmentation and selection using SVM hyper-plane is proposed in [14]. An audio track represented by long term statistical distribution of some short term spectral features (mel frequency cepstral coefficients (MFCC)) is used for ASC in [15]. Acoustic scene classification using matrix factorization with unsupervised feature learning is carried out in [16]. A compact and discriminative feature based on auditory summary statistics for ASC is proposed in [17]. Bag of Visual Words (BoVW) representations are used for ASC in [18]. Wavelets are revisited for the classification of acoustic scenes in [19,20]. Wavelet transform based mel-scaled features are used for ASC in [21]. Ensemble of spectrograms based on adaptive temporal divisions for ASC is used in [22]. Fully CNNs and I-Vectors were proposed in [23]. X-Vector embedding and CNN model for ASC were proposed in [24]. From the current research in the field of deep learning for all machine learning applications, and its success in DCASE 2016, DCASE 2017 [25] and DCASE 2018 [1].

Acoustic scenes datasets captured from the surrounding environments cover the audio frequency range from 20Hz to 20kHz. Features that can capture local information in both time and frequency domain would provide good representation of acoustic scenes. The results from previous research on DCASE 2013, 2016, 2017, and 2018 and analyses show the need for a suitable feature-classifier combination. Also, the consequent acoustic scenes constitute a variety of environmental sounds that can form complex sounds. Therefore, only one particular type of feature may not be sufficient to effectively and discriminatively represent them. In this paper, we propose the use of PLP and its variants for ASC using DNN as a classifier and also DNN score level fusion for decision.

The paper is organized as follows: Section 2 describes detail of feature extraction, Section 3 describes the proposed method. Section 4 discusses the results and analysis of the proposed system. Section 5 provides the summary and conclusions.

II. FEATURE EXTRACTION

In the present work, we are investigating PLP and PLPCC features for ASC. The description of these features are given below.

A. Perceptual Linear Prediction

The perceptual linear prediction model was developed by Hermansky [26]. PLP was proposed to match characteristics of human auditory system. Fig. 1. shows the essential components of PLP feature extraction. PLP has three important perceptual aspects: first- the critical-band resolution analysis, second- the equal-loudness analysis, and third- intensity-loudness conversion (cubic-root). Initially, the windowed signal power spectrum is calculated as,

$$P(\omega) = \text{Re}[S(\omega)]^2 + \text{Im}[S(\omega)]^2 \quad (1)$$

The first component is a conversion from audio frequency to Bark scale frequency, Bark frequency is a good representation of the human hearing resolution in frequency. The bark frequency is equivalent to an

Table 1. Acoustic scene classification results for DCASE 2018 task 1 subtask A (P1: PLP with DNN, P2: PLPCC with DNN, P3:

Log-Mel with DNN, P4: DNN score level fusion of PLP and PLPCC, P5: DNN score level fusion of PLPCC and Log-Mel, P6:

DNN score level fusion of PLP and Log-Mel, P7: DNN score level fusion of PLP, PLPCC and Log-Mel).

Acoustic Scene(%)	Baseline-2018[1]	P1	P2	P3	P4	P5	P6	P7
Airport	72.9	71.3	66.0	55.5	69.1	67.5	70.6	69.8
Bus	62.9	72.7	71.9	72.3	73.6	72.7	74.8	74.4
Metro	51.2	70.1	69.0	62.8	71.6	73.6	70.5	72.4
Metro station	55.4	65.6	65.3	63.3	66.8	65.3	65.3	64.5
Park	79.1	84.7	83.5	81.0	84.3	83.9	84.7	84.7
Public square	40.4	54.6	50.9	51.9	55.1	47.7	52.8	50.5
Shopping mall	49.6	49.1	59.9	83.5	52.3	78.9	74.9	77.8
Street pedestrian	50.0	57.9	59.5	52.6	58.3	56.3	55.5	55.9
Street traffic	80.5	89.0	89.0	89.4	89.0	89.0	88.6	88.6
Tram	55.0	71.3	68.2	70.5	69.3	72.0	72.0	72.0
Average	59.7(±0.7)	68.6	68.3	68.3	68.9	70.7	71.0	71.1

Table 2. Acoustic scene classification results for task 1 subtask B for the proposed system (P7).

Acoustic Scene(%)	Baseline-2018[1]				P7			
	A	B	C	Average(B,C)	A	B	C	Average(B,C)
Airport	73.4	68.9	76.1	72.5	61.9	27.8	44.4	36.1
Bus	56.7	70.6	86.1	78.3	71.1	83.3	83.3	83.3
Metro	46.6	23.9	17.2	20.6	66.3	38.9	72.2	55.5
Metro Station	52.9	33.9	31.7	32.8	61.4	38.9	38.9	38.9
Park	80.8	67.2	51.1	59.1	84.7	94.4	100.0	97.2
Public square	37.9	22.8	26.7	24.7	48.6	38.9	50.0	44.4
Shopping mall	46.4	58.3	63.9	61.1	74.6	88.9	77.8	83.3
Street pedestrian	55.5	16.7	25.0	20.8	62.3	33.3	44.4	38.8
Street traffic	82.5	69.4	63.3	66.3	86.2	77.8	83.3	80.5
Tram	56.5	18.9	20.6	19.7	74.7	38.9	44.4	41.6
Average	58.9(±0.8)	45.1(±3.6)	46.2(±4.2)	45.6(±3.6)	69.2	56.1	63.9	60.0

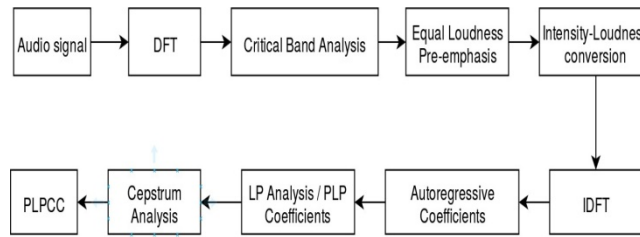


Fig. 1. Block diagram of the PLP cepstral coefficients feature extraction.

audio frequency is,

$$\Omega(\omega) = 6 \ln \omega / 1200\pi + [(\omega / 1200\pi)^2 + 1]^{0.5} \quad (2)$$

Second component, the auditory power spectrum is convoluted with the power spectrum of the critical-band masking curve to simulate the human hearing critical-band integration. Finally, smoothed power spectrum is downsampled at intervals of ≈ 1 Bark intervals. The three components frequency warping, smoothing and sampling are integrated into a one filter-bank called Bark filter bank. An equal-loudness pre-emphasis weighs filter-bank results to simulate the hearing sensitivity. The equalized values are changed according to the power law of Stevens by raising each to the power of 0.33.

The yielding auditory warped power spectrum is then processed by linear prediction (LP). Applying LP to the auditory warped power spectrum means, compute the predictor coefficients of a signal that has warped spectrum as a power spectrum and the predictor coefficients are termed as PLP coefficients. Finally, cepstral coefficients are achieved from the predictor coefficients by a recursion that is equivalent to the logarithm of the model spectrum followed by an inverse Discrete Fourier transform (IDFT). These coefficients are termed as PLP cepstral coefficients (PLPCC) [27].

The feature set is extracted for 40 filterbanks with a frame size of 40 ms with 20 ms overlap. Even for stacked context of features, static PLP (40 dimensions), Δ PLP (40 dimensions) and $\Delta\Delta$ PLP (40 dimensions) feature vectors are extracted making it 120 dimensions per frame. The same approach is followed for PLPCC feature extraction.

III. PROPOSED SYSTEM

The proposed ASC system incorporates the general ASC framework as shown in Fig. 2. From development data, all audio signals pass through pre-processing and feature extraction processes. Since the DCASE data is in binaural stereo format, the first pre-processing step is to convert the data sample to monophonic audio by averaging the two channels. Pre-emphasis factor of 0.97 was used to emphasize the high-frequency content. The features (PLP, PLPCC and Log-Mel band energies) are extracted from the preprocessed data. Models were built from features of training data and then employed for classification of the test samples using DNN model.

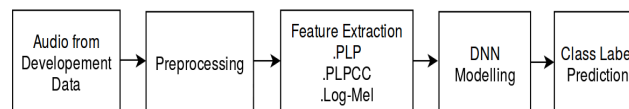


Fig. 2. Block diagram of proposed system for individual features.

A. Description of the Proposed System

In the proposed approach, we have considered DNN model as it will have fewer computations and also less training parameters than the baseline model CNN [1]. From this section onwards Log-Mel band energies is noted as Log-Mel. We have experimented with individual feature sets and combinations of DNN scores have been obtained with each feature set. The following are the seven set of proposed experiments:

P1: PLP with DNN, P2: PLPCC with DNN, P3: Log-Mel with DNN, P4: DNN score level fusion of PLP and PLPCC, P5: DNN score level fusion of PLPCC and Log-Mel, P6: DNN score level fusion of PLP and Log-Mel, P7: DNN score level fusion of PLP, PLPCC and Log-Mel

B. Detail of DNN Classifier

The DNN used in this study is a fully connected feedforward neural network. The network has an input layer, three hidden layers, and an output layer.

The input layer has 120 (dimension of feature vector) neurons with linear activation function while each hidden layer has 200 neurons with rectified linear unit (ReLU) activation function. The output layer has neurons (number of classes) with softmax activation and categorical cross-entropy loss function. ADAM [28] optimizer is used here for better weight optimization and its learning rate is set to 0.005. L2 regularizer is used to avoid overfitting with a value of 0.000099 and DNN trained with 30 epochs.

C. Decision Strategy

For individual feature sets, we used Max rule for classification. Computation of DNN score fusion of any two different features is done as follows: Let us consider the fusion of PLP and Log-Mel

band energies. If X_{PLP}^{i} and $X_{Log-Mel}^{i}$ are the DNN scores generated by two models for the ith acoustic scene, then a combined score is given by

$$X_d^{i, combine} = \alpha X_{PLP}^{i} + (1 - \alpha) X_{Log-Mel}^{i} \quad (3)$$

Similar procedure is carried out to compute DNN score fusion for different combinations of two feature sets. Computation of DNN score fusion of three different features is given by

$$X_d^{i, combine} = \alpha X_{PLP}^{i} + \beta X_{PLPCC}^{i} + \gamma X_{Log-Mel}^{i} \quad (4)$$

where, summation of α, β and γ is one.

In these proposed systems, for P1 to P6, $\alpha = 0.5$ and for P7, $\alpha = 0.5, \beta = 0.2$ and $\gamma = 0.3$ is fixed to obtain significant improvements.

where, summation of α, β and γ is one.

In these proposed systems, for P1 to P6, $\alpha = 0.5$ and for P7,

$\alpha = 0.5, \beta = 0.2$ and $\gamma = 0.3$ is fixed to obtain significant improvements.

IV. RESULTS AND DISCUSSIONS

This section gives the datasets used, baseline methods, results and discussion on ASC tasks for DCASE 2018 Task 1 subtask A (basic ASC) and subtask B (ASC with mismatched recording devices) and DCASE 2017 (ASC). The results of DCASE 2018 task 1 subtask A are presented in Table 1 and results of subtask B are presented in Tables 2 and 3 respectively. The results of DCASE 2017 ASC Task are presented in Tables 4 and 5.

A. Datasets and Baseline Methods

We have used the development dataset of TUT Acoustic Scenes 2018 [1] and TUT Acoustic Scenes 2017 [25]. According to the DCASE 2017 challenges' ASC task setup, development data is partitioned into k folds, where k=4 for both the datasets. Fold-wise mean classification accuracy is used as the performance metric during development. For performance comparison, we have used the baseline systems of the DCASE challenges of 2018 and 2017, which are Log-Mel band energies with CNN [1] and Log-Mel band energies with MLP [25].

B. Analysis of Results on DCASE 2018 Subtask A

The results for the subtask A are given in Table 1 for the DCASE 2018 baseline system and proposed systems (P1-P7). This subtask is concerned with the basic problem of ASC, in which all available data is recorded with the same device, which in this case is device A (Zoom F8 audio recorder).

From the table, it can be observed that individual PLP features perform better than PLPCC and Log-Mel features. For two features combination proposed systems (P4-P6), performance has increased consistently, out of which P6 has given better performance than P4 and P5 with a relative performance of 11.3% as compared to baseline system.

Further, it is observed that the DNN score fusion of PLP, PLPCC and Log-Mel band energies features (P7) has given significant improvement in accuracy and indicates complementary information based on 11.4% of relative improvement. Using proposed features (P7) the Shopping mall, Park and Street traffic classes are well classified when compared to other classes for subtask A. From the table, it can also be observed that the proposed system has given an improvement in the average accuracy. The relative improvements of 8.9%, 8.6%, 8.6%, 9.2%, 11.0%, 11.3%, and 11.4% are obtained for P1-P7 respectively as compared to the DCASE 2018 baseline system.

C. Analysis of Results on DCASE 2018 Subtask B

The results for the subtask B are presented in Tables 2 and 3. This subtask is concerned with the situation in which an application will be tested with a few different types of devices (device A, device B- Samsung Galaxy S7 and device C- iPhone SE), preferably not the same device as the ones used to record the development data. The results of the Subtask B are given in Table 2 for the DCASE 2018 baseline and the proposed system which uses the DNN score fusion of PLP, PLPCC and Log-Mel band energies features (P7). From this table, it can be noted that individual PLP features perform better than PLPCC and Log-Mel features. For two features combination proposed systems (P4-P6), performance has increased consistently, out of which P6 gave better performance than P4 and P5 with a relative performance of 14.2% compared to DCASE 2018 baseline system. It can be observed that the DNN score fusion of PLP, PLPCC and Log-Mel band energies feature in the proposed system (P7) has given significant improvement compared to the DCASE 2018 baseline system. Overall, 14.4% relative improvement is achieved with the proposed system. Using proposed features (P7) the Bus, Park, Street traffic and Shopping mall classes are well

Table 3. Average (B, C) accuracies for Subtask B.

Accuracy	Baseline-2018[1]	P1	P2	P3	P4	P5	P6	P7
Average	45.6%(±3.6)	59.2%	58.1%	54.5%	59.3%	59.6%	59.8%	60.0%

Table 4. Average accuracies for DCASE 2017 dataset on 4 fold cross validation.

Accuracy	Baseline-2017[25]	P1	P2	P3	P4	P5	P6	P7
Average	74.8	80.6%	79.5%	75.9%	80.1%	80.4%	80.6%	80.8%

classified when compared to other classes for subtask B. Further, we have also experimented with the other proposed systems (P1, P2, P3, P4, P5, P6, and P7). The average (B, C) performance was obtained with various proposed feature sets shown in Table 3. From the table, it can be observed that all the proposed systems have given an improvement in the average (B, C) accuracy. The relative improvements of 13.6%, 12.5%, 8.9%, 13.7%, 14.0%, 14.2%, and 14.4% are obtained for P1-P7 respectively as compared to the DCASE 2018 baseline system.

D. Analysis of Results on TUT Acoustic Scenes 2017 Dataset

The results on TUT Acoustic Scenes 2017 dataset for AS task are presented in Tables 4 and 5.

From the Table 4, individual PLP features perform better than PLPCC and Log-Mel features. For two features combination proposed systems (P4-P6), performance has increased consistently, out of which P6 has given better performance than P4 and P5 with a relative performance of 5.8% compared to DCASE 2017 baseline system. It can be observed that all the proposed systems have given an improvement in the average accuracy. The relative improvements of 5.8%, 4.7%, 1.1%, 5.3%, 5.6%, 5.8%, and 6.0% are obtained for P1-P7 respectively as compared to the DCASE 2017 baseline system.

The results on TUT Acoustic Scenes 2017 dataset are given in Table 5 for the DCASE 2017 baseline and our proposed system, which is the DNN score fusion of PLP, PLPCC and Log-Mel band energies features (P7) and also the comparison with state-of-art in [29].

Table 5. ASC results, averaged over evaluation folds and comparison with DCASE 2017 Baseline system and state of the art [14].

Acoustic Scene	Baseline-2017[25](%)	Fusion without augmentation[14]	Proposed System (P7)(%)
Beach	75.3	70.9	51.3
Bus	71.8	82.1	87.2

Cafe/Restau- rant	57.7	71.8	64.1
Car	97.1	89.0	100.0
CityCenter	90.7	85.6	94.9
ForestPath	79.5	97.3	47.4
GroceryStor- e	58.7	83.3	89.7
Home	68.6	76.0	89.7
Library	57.1	82.0	80.8
MetroStation	91.7	90.7	98.7
Office	99.7	95.1	97.4
Park	70.2	69.9	55.1
ResidentialA- rea	64.1	71.8	94.9
Train	58.0	71.8	69.2
Tram	81.7	84.6	91.0
Overall accur- acy	74.8	81.5	80.8

It can be seen from the table that the proposed features (P7) are well classified for Car, City center and Office classes compared to DCASE 2017 baseline [25] and [29].

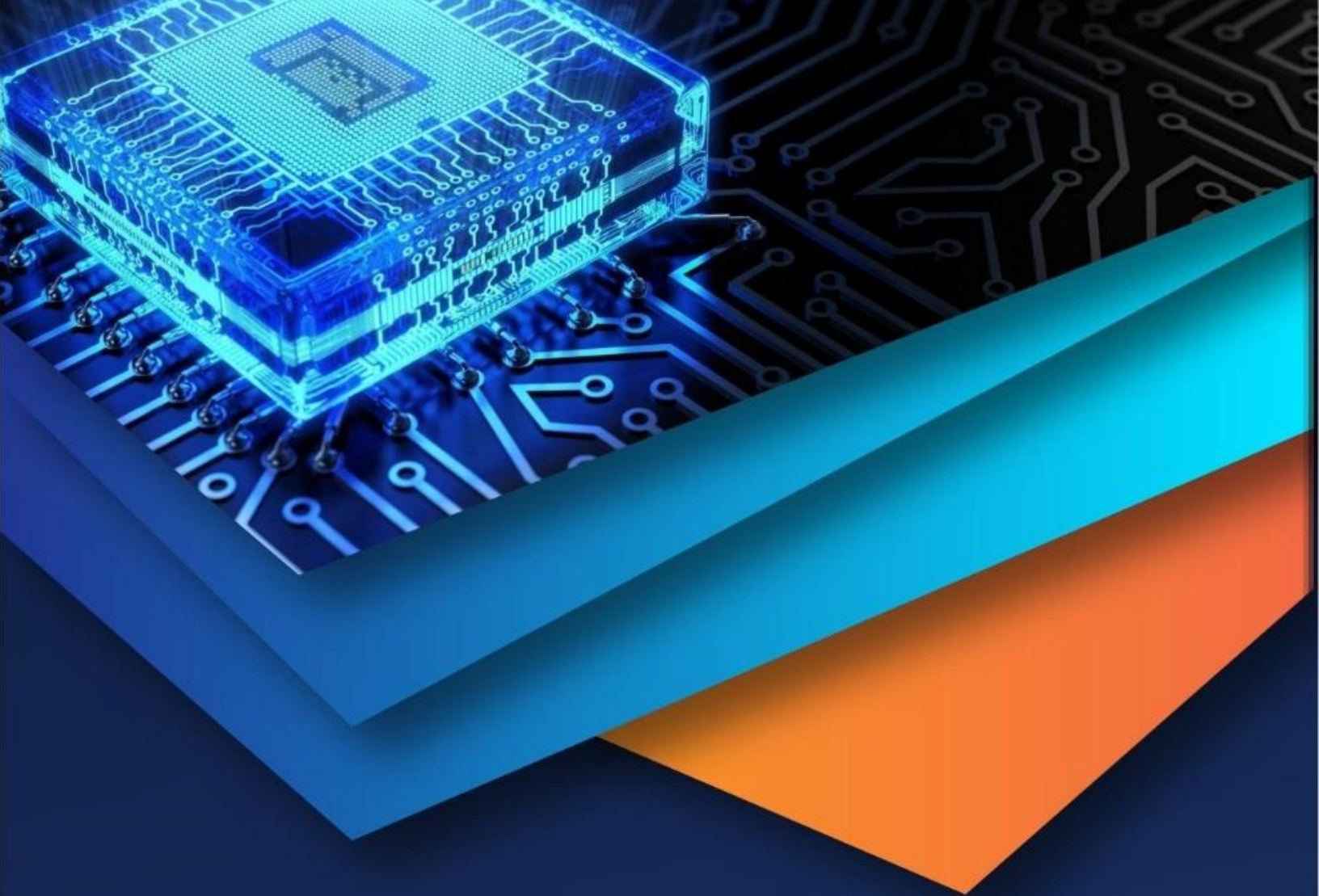
V. SUMMARY AND CONCLUSIONS

In this paper, an investigation of PLP and PLPCC features with DNN architecture has been applied to model the ASC. We experimented with TUT Acoustic Scenes 2018 Datasets of task1 including Sub-task A and B and TUT Acoustic Scenes 2017 dataset. The study demonstrated that the capability of individual feature sets and fusion of PLP, PLPCC and Log-Mel band energies at DNN score decision level. Individual PLP features yield an improvement of 8.9% and 13.6% and PLPCC features result in an improvement of 8.6% and 12.5% for subtask A and subtask B of DCASE 2018 challenge. Significant improvements in accuracy is achieved for DNN decision levels scores obtained from PLP, PLPCC and Log-Mel band energies. Improvement of 11.4% and 14.4% were achieved in subtasks A and B respectively compared to the DCASE 2018 ASC baseline system. From DCASE TUT Acoustic Scenes 2017 dataset, individual PLP features yield an improvement of 5.8% and PLPCC features result in an improvement of 4.7% respectively. An improvement of 6.0% is achieved with the fusion study compared to the DCASE 2017 baseline system. This shows that PLP, PLPCC and Log-Mel band energies carry complementary acoustic information. Future work would be dedicated to the investigation of different combinations of features for ASC.

REFERENCES

- [1] Mesaros, Annamaria and Heittola, Toni and Virtanen, Tuo- mas, "A multi-device dataset for urban acoustic scene clas- sification," Submitted to DCASE2018 Workshop, 2018.
- [2] Daniele Barchiesi, Dimitrios Giannoulis, Dan Stowell, and Mark D. Plumbley, "Acoustic scenes classification: Classifying environments from the sounds they produce," The Journal of IEEE Signal Processing Magazine, vol. 32, pp. 16–34, 2015.
- [3] S. Chu, S. Narayanan, C. J. Kuo, and M. J. Mataric, "Where am I? scene recognition for mobile robots using audio features," in Proc. IEEE International Conference on Multimedia and Expo, 2006, pp. 885–888.
- [4] J. T. Geiger, B. Schuller, and G. Rigoll, "Recognising acoustic scenes with large-scale audio feature extraction and SVM," in Proc. IEEE AASP Challenge on Detection and Classification of Acoustic Scenes and Events, 2013.
- [5] W. Nogueira, G. Roma, and P. Herrera, "Sound scene identi- fication based on MFCC, binaural features and a support vector machine classifier," in Proc. IEEE AASP Challenge on Detec- tion and Classification of Acoustic Scenes and Events, 2013.
- [6] A. Rakotomamonjy and G. Gasso, "Histogram of gradient of time-frequency representations for audio scene detection," IEEE Transactions on Audio, Speech, and Language Process- ing, vol. 23, no. 1, 2015.
- [7] Axel Plinge, Rene Grzeszick, and Gernot A. Fink, "A Bag-of- features approach to acoustic event detection," in Proc. IEEE International Conference on Acoustics, Speech and Signal Pro- cessing (ICASSP), 2014, pp. 3704–3708.
- [8] P. Khunarsal, C. Lunsin, and T. Raicharoen, "Very short time environmental sound classification based on spectrogram pat- tern matching," The Journal of the Information Sciences, vol. 243, pp. 57–74, 2013.

- [9] JSalamonandJPBello, "Deepconvolutionalneuralnetworks and data augmentation for environmental sound classification," vol. 24, no. 3, pp. 279–283, 2017.
- [10] W Nogueira, G Roma, and P Herrera, "Sound scene identificationbasedonMFCC,binauralfeaturesandasupportvector machineclassifier," in Proc. IEEE AASP Challenge on Detection and Classification of Acoustic Scenes and Events, 2013.
- [11] M Chum, A Habshush, A Rahman, and C Sang, "Scene classification challenge using hidden Markov models and frame based classification," in Proc. IEEE AASP Challenge on Detection and Classification of Acoustic Scenes and Events, 2013.
- [12] HamidEghbal-Zadeh, BernhardLehner, MatthiasDorfer, and Gerhard Widmer, "CP-JKU submissions for DCASE-2016: a hybrid approach using binaural i-vectors and deep convolutionalneuralnetworks," Tech. Rep., DCASE2016Challenge, September 2016.
- [13] Rakib Hyder, Shabnam Ghaffarzadegan, Zhe Feng, John H L Hansen, and Taufiq Hasan, "Acoustic scene classification using a CNN-Supervectors system trained with auditory and spectrogram image features," in Proc. INTERSPEECH, 2017.
- [14] SeongkyuMun, SangwookPark, DavidHan, and HanseokKo, "Generative adversarial network based acoustic scene training set augmentation and selection using SVM hyper-plane," Tech. Rep., DCASE2017 Challenge, September 2017.
- [15] YPetetin, CLaroche, and AMayoue, "Deepneuralnetworks for audio scene recognition," in Proc. European Conference on Signal Processing (EUSIPCO), 2015, pp. 125–129.
- [16] V Bisot, R Serizel, S Essid, and G Richard, "Acoustic scene classification with matrix factorization for unsupervised feature learning," in Proc. IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), March 2016, pp. 6445–6449.
- [17] HongweiSong and Jiqing Han and ShiwenDeng, "A compact and discriminative feature based on auditory summary statistics for acoustic scene classification," in Proc. Interspeech, 2018, pp. 3294–3298.
- [18] Manjunath Mulimani and Shashidhar G Koolagudi, "Robust acoustic event classification using bag-of-visual-words," in Proc. Interspeech, 2018, pp. 3319–3322.
- [19] Qian, Kun and Ren, Zhao and Pandit, Vedhas and Yang, Zi-jiang and Zhang, Zixing and Schuller, Björn, "Wavelets revisited for the classification of acoustic scenes," in Proc. Workshop on Detection and Classification of Acoustic Scenes and Events (DCASE), November 2017, pp. 108–112.
- [20] Shefali Waldekar and Goutam Saha, "Wavelet-based audio features for acoustic scene classification," Tech. Rep., DCASE Challenge, September 2018.
- [21] Shefali Waldekar and Goutam Saha, "Wavelet transform based mel-scaled features for acoustic scene classification," in Proc. Interspeech, 2018, pp. 3323–3327.
- [22] Sakashita, Yuma and Aono, Masaki, "Acoustic scene classification by ensemble of spectrograms based on adaptive temporal divisions," Tech. Rep., DCASE Challenge, September 2018.
- [23] Dorfer, Matthias and Lehner, Bernhard and Eghbal-zadeh, Hamid and Christop, Heindl and Fabian, Paischer and Gerhard, Widmer, "Acoustic scene classification with fully convolutional neural networks and I-vectors," Tech. Rep., DCASE Challenge, September 2018.
- [24] Zeinali, Hossein and Burget, Lukas and Cernocky, Honza, "Convolutional neural networks and x-vector embedding for dcase2018 acoustic scene classification challenge," Tech. Rep., DCASE Challenge, September 2018.
- [25] Annamaria Mesaros, Toni Heittola, and Tuomas Virtanen, "TUT database for acoustic scene classification and sound event detection," in Proc. 24th Conference on European Signal Processing (EUSIPCO), Budapest, Hungary, 2016.
- [26] Hynek Hermansky, "Perceptual linear predictive (PLP) analysis of speech," The Journal of the Acoustical Society of America, vol. 87, no. 4, pp. 1738–1752, 1990.
- [27] Hacker Christian Hoenig Florian, Stemmer Georg and Brugnara Fabio, "Revising Perceptual Linear Prediction (PLP)," in Proc. INTERSPEECH, 2005, pp. 2997–3000.
- [28] J Schmidhuber, "Deep learning in neural networks: An overview," arXiv, 2014.
- [29] Diederik, P Kingma, and Jimmy Lei Ba, "ADAM: A Method for Stochastic Optimization," in Proc. Conference on ICLR, 2015.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)