

Applying convolutional neural networks to the analysis of mouse ultrasonic vocalizations

Reyhaneh ABBASI¹; Peter BALAZS¹; Anton NOLL¹; Doris NICOLAKIS²; Maria Adelaide MARCONI²; Sarah M. ZALA²; Dustin J. PENN²

¹ Acoustics Research Institute, Austrian Academy of Sciences, Austria

² Konrad Lorenz Institute of Ethology, University of Veterinary Medicine, Austria

ABSTRACT

Many species in diverse taxonomic groups, including rodents, bats, and insects, communicate with complex ultrasonic vocalizations (USVs) (>20 kHz). Two main components of processing and analyzing USV recordings include detection and classification of syllable types. Recently we developed an efficient algorithm for detecting mouse USVs (Automatic Mouse Ultrasound Detector (A-MUD)). The main challenge is detecting USVs under conditions with a low signal-to-noise ratio, while minimizing rates of false positives (FP). Mice produce many short USVs (< 10 millisecond) that inflate FPs. We aimed to improve the detection of mouse USVs with A-MUD by classifying vocalizations into three discrete syllable types (with 0, 1, or ≥ 2 frequency-jumps) or FP. Supervised Convolutional Neural Networks (CNNs) were fed by 2D Gammatone Spectrograms (GSs) adapted to the frequency range of mice. Evaluation of performance shows that CNNs yielded an overall accuracy of $95 \pm 1.2\%$ and macro-F1 score of $90 \pm 2.7\%$. In contrast, multilayer feed-forward neural networks fed by vectorized spectrograms provided an overall accuracy of only $85.4 \pm 1.9\%$ and macro-F1 score of $75.4 \pm 2.9\%$, and therefore, the chosen CNNs outperformed this conventional classification method.

Keywords: Deep learning, Ultrasonic vocalizations (USVs), Image-based classification

1. INTRODUCTION

Mice like other taxonomic groups use ultrasonic vocalizations (USVs) (20-120 kHz) to communicate. USVs are emitted by mouse pups when they are in distress or away from their mother, and by adults when they socially interact (1, 2). The function of USVs in house mice may be to convey information about a signaler's identity (species, sex, age), sexual motivation, health and social status (3-7). USVs are increasingly used for investigating the genetic basis of autism and human speech disorders (8, 9). Analyzing USV data is extremely time-consuming, however, and better methods for automatic detection and classification of USVs are needed.

There are commercial and non-commercial software tools for automatic detection of USVs available, including the following: MUPET (10), mouse song analyzer (11), Avisoft SASLab Pro (Version 4.2, Avisoft Bioacoustics, Berlin, Germany) and SONOTRACK (Version 2.2.4, Metris, Netherlands). Their performance is acceptable for recordings of mice having a high signal-to-noise ratio (SNR), but noisy recordings, such as vocalizations emitted during sexual or social interactions, are problematic. We developed an automatic method, the mouse ultrasound detector (A-MUD) (version 1.0), which has very good performance under such conditions (12). Our most recent version (3.2) (13) has an even lower false negative rate (FNR), caused by missing short syllables (≤ 10 milliseconds) and faint (low amplitude) USVs; however, it has a higher false positive rate (FPR). This is an inescapable tradeoff for any signal detection method, as noted by others (14), though we nevertheless aim to improve the performance.

In the development of previous methods, the removal of FPs was conducted both manually (15, 16) and automatically (10, 14, 17). Automatic elimination of FPs can be accomplished with unsupervised and supervised approaches. Van Segbroeck et al. (10) implemented a K-means clustering algorithm (unsupervised approach) (18) in MUPET. They classified detected elements (i.e., USVs or FPs) into

¹ reyhaneh.abbasi@oeaw.ac.at

100 classes and the user must select which class(es) represent FPs to be removed. K-means clustering allows low computational load (10), while it still leaves some FPs unidentified. Coffey et al. (14) also attempted to eliminate FPs by supervised trained convolutional neural networks (CNNs). Using a deep learning-based approach (19), spectrogram images have been first classified into two classes, FPs versus USVs, and then, the identified USVs are reclassified. Although this method has an acceptable level of accuracy, the stepwise elimination of FPs and classification USVs is time-consuming (20). Moreover, this method assumes that no FPs remain in the data after the first step, and it therefore considers no target class for FP in the second step. Our studies show that this assumption not always holds, and it may degrade the performance of the model.

Methods for automatic classification of USV syllable types also employ both unsupervised (10, 21) and supervised (14) methods. In the supervised classification (22, 23), observations are manually labeled and represented as response variable to the model. Higher accuracy is secured when using this method rather than (unsupervised) clustering (24, 25). To the best of our knowledge, DeepSqueak is the only available software package that provides supervised classification of USVs using CNNs. The model is trained on a relatively high number of samples (56,000 USVs) acquired from the Mouse Tube dataset. This model categorizes inputs based on spectrogram images into 5 default classes including Split, Inverted U, Short Rise, Wave, and Step. Our examinations of DeepSqueak's performance revealed some deficiencies. First, as previously mentioned, FPs are not assigned with a target class, which may increase errors in classification. The selection of target classes is normally based on the importance of these classes in the behavioral studies of mice (25, 26). However, the authors of this method have not discussed the reason for choosing the aforementioned classes. In addition, the performance of DeepSqueak classification has not been compared with other models. Finally, CNNs used in DeepSqueak is developed on the high-SNR Mouse Tube data. Therefore, a full assessment of DeepSqueak is postponed until its application to high-noise data.

The aims of the present study include the following: (1) classify USVs into four classes, FP, C2, C3, and no-jump (NJ), as the most considered classes in mice behavior studies, (2) use low-SNR data for training CNNs, and (3) evaluate the feasibility of using spectrogram images as the input for the classifier. The latter is considered as the replacement for vectorized images.

2. DATA and METHOD

2.1 Data

Here, we used USV recordings of adult wild-derived house mice (*Mus musculus musculus*). Animal housing and weaning conditions have been previously described (12). A-MUD (version 3.2) is used for the detection of elements. Then, they were manually labeled into 4 classes: FP, C2 (one frequency jump), C3 (two or more jumps) and NJ (no jumps) (26, 27). For this study, we arbitrarily selected 1000 members of the class NJ (out of 7000 members), 120 members from C2 and FP, and 75 members from C3.

2.2 Method

2.2.1 Preprocessing

For image-based classification, data is given in the form of two-dimensional images (whether being an image of a flower or a 2D spectrogram) to a classifier. In the present study, these images were derived from the spectrograms of detected elements. Spectrogram computation consists of two steps: filtering and applying the short time Fourier transform (STFT) (28). Initially, the signal is passed through Chebyshev filter (29) (Type I, order 8, bandpass with corners of $F_{min} = 30$ kHz, and $F_{max} = 120$ kHz) to extract the desired signal in the ultrasound range. In the second step, according to parameters already defined by A-MUD, spectrogram is produced via STFT. Having prepared the mice USVs spectrograms, they are postprocessed in two steps. First, for the sake of computational expense, following Van Segbroeck et al. (10), the Gammatone filters (30) are used to reduce the size of spectrogram array along the frequency axis. Because this filter is primarily defined for the human perception of sound (31), here we adapted it to the frequency range of mice hearing using the filter specifications defined in MUPET. Secondly, the function maximum is applied to filtered spectrogram and floor noise (10^{-3}). The output is logarithmically transformed and, then, smoothed by using auto regression moving-average filter (32) with order 1. The resulting smoothed spectrogram has a size of 64×401 (compare to the original size of 225×401) which is the number of frequency bins multiplied by number of time bins. Hereafter, it is called Gammatone spectrograms (GSs).

2.2.1 Classifier

Here, we have compared the performance of two classifiers including CNNs and multilayer feed-forward neural networks (MFNNs). Although both techniques are categorized as being deep learning algorithm, the main difference is that the former is fed by 2D images while the latter normally uses vectorized images as input. Conceptually to preserve the relation of time and frequency, the first approach seems far more feasible in our setting. CNNs was first introduced by Fukushima (33) and developed by LeCun et al. (34). The successful performance of CNNs in categorizing images has been already documented in various fields of study including handwritten recognition (35), object discrimination (36), categorization of birds species (37) and identification of environmental sound sources (38). MFNNs also performed well in audio analysis disciplines such as marmosets' vocalizations' classification (39), sound event classification (40), and speech recognition (41). This study investigates the use of CNNs, fed by 2D images, for improving USVs analysis. It is inspired by the fact that MFNNs with vectorized inputs may lead to a poor classification (42).

The architecture of MFNNs consists of 4 FC layers (each with 50 hidden neurons), each followed by a dropout layer (43) with a probability of 0.5, and the output FC layer with 4 neurons.

For the development of CNNs, we have used the following structure: the input layer, 8 convolutional layers (3*3) for extraction of features from images, each followed by an activation function of rectified linear unit (ReLU) (44) (set by default parameters), two non-consecutive layers of max pooling (2*2) (45) for reduce the size of the input, and two FC layers for classification. The architecture of both models is selected based on trial and error.

To optimize the loss function of categorical cross-entropy in CNNs and MFNNs, Adam algorithm is used (46). This loss function computes the dissimilarity between the distribution of predictions (after applying the non-linear softmax function) and of the true data using the following equation:

$$\text{Categorical cross entropy} = \sum_{i=1}^C t_i \log(f(s_i)), \quad (1)$$

Where t_i and s_i are the true and estimated distributions of data, f is the softmax function $e^{s_i} / \sum_j^C e^{s_j}$ and C is number of classes. In order to prevent overfitting, augmentation (36) is applied to the training data. To solve the problem of imbalanced number of class members, we used the incorporation of classes weight in the loss function (47). These weights are obtained by the inverse of the number of class members. Finally, 80% of the observations (per class) and their respective features are used to train the classifiers. The classification is performed on 5 folds of dataset.

3. RESULTS

In this section we will discuss the class separability using unsupervised t-distributed stochastic neighbor embedding (t-SNE) (48). Then, we will present the performance analysis of CNNs and MFNNs.

3.1 Specifications of Classes

In Figure 1-a, GSs of 5 randomly selected members from each studied class is shown. The classes C2 and C3 are both composed of 2 or more notes (defined as basic acoustic unit which are formed by a single continuous sound with gradual variations in fundamental frequency (49)). GSs of the class FP only contains noise. Members of NJ (row 3) do not contain jumps, but the overall shape of some examples is similar to some jump-contained USVs or the class FP. Hence, the 3D visualization of t-SNE (Figure 1-b) shows FP and C2 members overlapped NJ members. But the members of the classes C2, C3, and FP are located in a distinguishable distance. It is necessary to note that there is the possibility of the co-occurrence of each of the above USV classes members with background noise (row 2 column 1 and 4), or USVs might be faint (row 3 column 1), which in this study is not considered as a distinct class (e.g. class of noisy USVs).

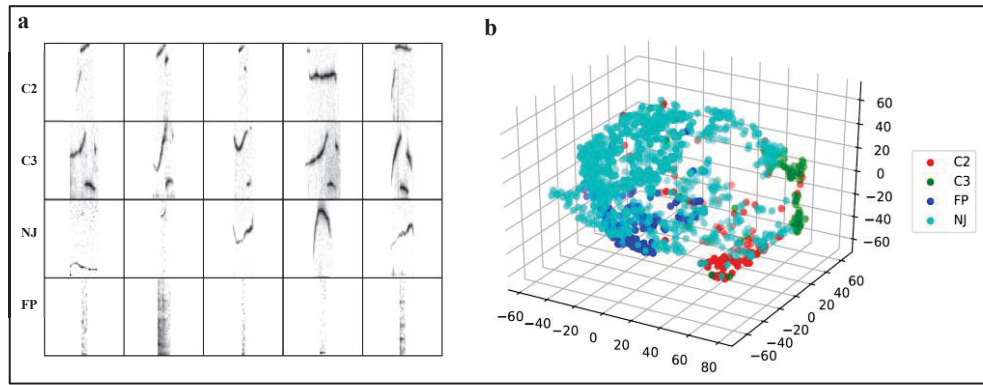


Figure 1 – a) GSs of 5 randomly selected members of four studied classes and b) 3D visualization of t-SNE values. Classes are represented by color.

3.2 Classifiers Performance

The following table summarizes the performance of CNNs and MFNNs (derived from 5 folds) in classifying detected elements.

Table 1 – The performance of CNNs and MFNNs in classifying detected elements into 4 classes. The values are means $\pm$ one standard deviation. Please note that overall metrics are unweighted average of classwise metrics.

Models	Overall metrics			Area under the ROC curve (AUC) (50) (classwise)			
	Accuracy (%)	F1 (%)	AUC	C2	C3	NJ	FP
CNNs	95 $\pm$ 1.2	90 $\pm$ 2.7	0.94 $\pm$ 0.01	0.91 $\pm$ 0.04	0.94 $\pm$ 0.03	0.94 $\pm$ 0.01	0.95 $\pm$ 0.02
MFNNs	85.4 $\pm$ 1.9	75.4 $\pm$ 2.9	0.81 $\pm$ 0.05	0.8 $\pm$ 0.06	0.77 $\pm$ 0.1	0.81 $\pm$ 0.04	0.85 $\pm$ 0.03

According to the standard metrics of accuracy, macro-F1 score, and the area under the ROC curve (AUC) (50), CNNs gives much better results than MFNNs. Moreover, the variation of MFNNs performance metrics between classes is much higher than that of CNNs. For example, MFNNs resulted in the AUC of 0.85 $\pm$ 0.03 in the class FP and 0.77 $\pm$ 0.1 in the class C3, whereas the CNNs equals to 0.95 $\pm$ 0.02 and 0.94 $\pm$ 0.03, respectively. Considering the outperformance of CNNs, the rest of the paper is allotted to the detail evaluation of this algorithm.

In Figure 2, CNNs performance is evaluated based on ROC (Figure 2– a) and normalized confusion matrix (Figure 2– b). In addition, cases that are erroneously classified by CNNs are depicted in Figure 2– c. In Figure 2 – a, the ROC curves are accompanied with the AUC values in the legend. The higher AUC quantities (ROC closer to top left) show the greater ability of the model to separate the corresponding class from the others. According to this figure, the classes FP and C2 have the highest and the lowest separability, respectively. The reason for this can be found in Figure 1 – b, where the class FP (unlike the class C2) is located in high distance from other classes. In order to better understand the CNNs performance, the FPR of the model (those members off the diagonal axis) for each class is presented in Figure 2 – b. In this figure, the use of normalization (by rows) is used to eliminate the misleading effects of imbalanced classes on the model's performance metrics. Accordingly, the maximum within-group error refers to the misclassification of the class C2. Members of this class are mainly mistakenly predicted as C3 and NJ. The best result is obtained for the class FP with a prediction accuracy of 100%. These results are most likely not affected by the number of class members. As the number of members is equal in three classes FPs, C2 and C3 for train and test.

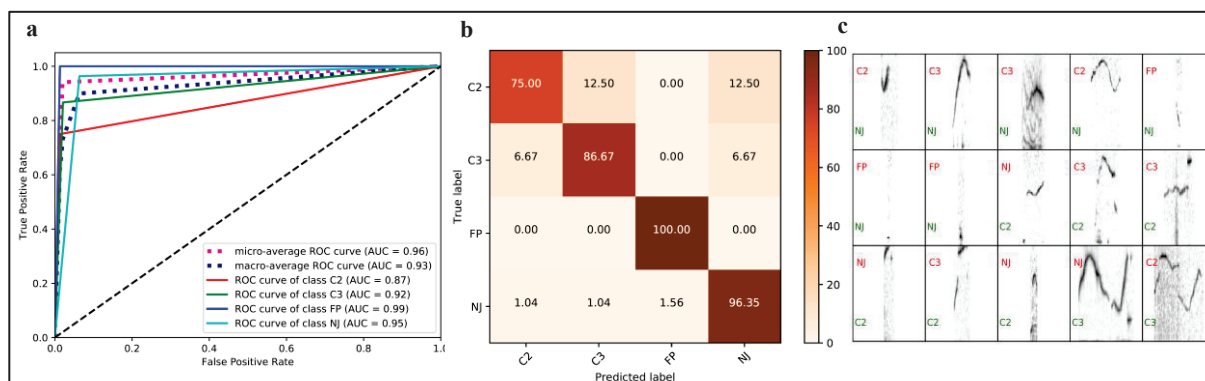


Figure 2 – a) ROC curve for each class and micro (weighted)- and macro- (unweighted) average of all classes. Micro average is weighted by the number of members. b) Normalized (by row) confusion matrix of CNNs, and c) all elements whose labels are mistakenly predicted by CNNs. The green and red labels represent the correct and false labels, respectively.

In Figure 2 – c, examples of falsely classified members with their correct labels are shown. As expected, the NJ USVs, which mistakenly labeled as FP, have very close pattern to them. The probable cause of errors in identifying members of the classes NJ and C2 (which are classified as C3 or C2) can be attributed to the background noise. Finally, the classification error in the class C3 can be caused by the significant difference between GSs shown in Figure 2 – c (row 3 column 4) and the representative pattern of this class.

4. CONCLUSION

In this study, we used the image-based CNNs to improve the classification of detected elements from mice audio recordings. Target classes, which were selected based on their importance in the studies of mice behavior, are three classes of C2, C3, NJ and FP. The comparison of this method with MFNNs, fed by vectorized images, shows a much higher efficiency of CNNs. Further examinations of CNNs classification show that the background noise in USVs causes the misclassification of the classes NJ and C2. Additionally, error rate increases in cases where the pattern in a segment is different from the representative pattern of a class. This problem has been identified as the main origin of errors in the class C3. We suspect that the performance of CNNs can be improved by using more data, more efficient noise removal methods and implementing other CNNs architectures. In the future, we will compare the results of this model with DeepSqueak, as the state-of-the-art model for USVs classification.

ETHICAL STATEMENT

This study was carried out in strict accordance with the recommendations in the Guide for the Care and Use of Laboratory Animals of the National Institutes of Health. All the experiments were conducted at the Konrad Lorenz Institute of Ethology, Austria and the protocols have been approved and were in accordance with ethical standards and guidelines in the care and use of experimental animals of the Ethical and Animal Welfare Commission of the University of Veterinary Medicine, Vienna (Austria). We did not sacrifice any of the mice used for this study.

ACKNOWLEDGEMENTS

This work was supported by Austrian Science Fund (AT) P 28141-B25 to DJP and SMZ, and the START-project FLAME ('Frames and Linear Operators for Acoustical Modeling and Parameter Estimation'; Y 551-N13) of PB.

REFERENCES

1. Sales G, Pye D. Ultrasonic communication by animals. London: Chapman & Hall; 1974.
2. Knutson B, Burgdorf J, Panksepp J. Ultrasonic vocalizations as indices of affective states in rats. *Psychological Bulletin*. 2002;128(6):961.
3. Kobrina A, Dent ML. The effects of aging and sex on detection of ultrasonic vocalizations by adult CBA/CaJ mice (*Mus musculus*). *Hearing Research*. 2016;341:119-29.
4. Von Merten S, Hoier S, Pfeifle C, Tautz D. A role for ultrasonic vocalisation in social communication and divergence of natural populations of the house mouse (*Mus musculus domesticus*). *Plos One*. 2014;9(5):e97244.
5. Chabout J, Serreau P, Ey E, Bellier L, Aubin T, Bourgeron T, et al. Adult male mice emit context-specific ultrasonic vocalizations that are modulated by prior isolation or group rearing environment. *PloS One*. 2012;7(1):e29401.
6. Zampieri BL, Fernandez F, Pearson JN, Stasko MR, Costa AC. Ultrasonic vocalizations during male–female interaction in the mouse model of Down syndrome Ts65Dn. *Physiology & Behavior*. 2014;128:119-25.
7. Zala SM, Reitschmidt D, Noll A, Balazs P, Penn DJ. Sex-dependent modulation of ultrasonic vocalizations in house mice (*Mus musculus musculus*). *PloS One*. 2017;12(12):e0188647.
8. Fischer J, Hammerschmidt K. Ultrasonic vocalizations in mouse models for speech and socio - cognitive disorders: insights into the evolution of vocal communication. *Genes, Brain and Behavior*. 2011;10(1):17-27.
9. Scattoni ML, Gandhi SU, Ricceri L, Crawley JN. Unusual repertoire of vocalizations in the BTBR T+tf/J mouse model of autism. *PloS One*. 2008;3(8):e3067.
10. Van Segbroeck M, Knoll AT, Levitt P, Narayanan S. MUPET—Mouse Ultrasonic Profile ExTraction: A Signal Processing Tool for Rapid and Unsupervised Analysis of Ultrasonic Vocalizations. *Neuron*. 2017;94(3):465-85. e5.
11. Chabout J, Jones-Macopson J, Jarvis ED. Eliciting and Analyzing Male Mouse Ultrasonic Vocalization (USV) Songs. *Journal of Visualized Experiments: JoVE*. 2017(123).
12. Zala SM, Reitschmidt D, Noll A, Balazs P, Penn DJ. Automatic mouse ultrasound detector (A-MUD): A new tool for processing rodent vocalizations. *PloS One*. 2017;12(7):e0181200.
13. Zala SM, Nicolakis D, Marconi MA, Noll A, Balazs P, Penn DJ. Primed to sing: wild-derived male house mice alter the quantity and quality of their vocalizations after interacting with a female. Under review. 2019.
14. Coffey KR, Marx RG, Neumaier JF. DeepSqueak: a deep learning-based system for detection and analysis of ultrasonic vocalizations. *Neuropsychopharmacology*. 2019:1.
15. Peters SM, Tuffnell JA, Pinter IJ, van der Harst JE, Spruijt BM. Short-and long-term behavioral analysis of social interaction, ultrasonic vocalizations and social motivation in a chronic phencyclidine model. *Behavioural Brain Research*. 2017;325:34-43.
16. Wright JM, Gourdon JC, Clarke PB. Identification of multiple call categories within the rich repertoire of adult rat 50-kHz ultrasonic vocalizations: effects of amphetamine and social context. *Psychopharmacology*. 2010;211(1):1-13.
17. Smith AA, Kristensen D, editors. Deep learning to extract laboratory mouse ultrasonic vocalizations from scalograms. *IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*; 2017: IEEE.
18. MacQueen J, editor Some methods for classification and analysis of multivariate observations. *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*; 1967: Oakland, CA, USA.
19. LeCun Y, Bengio Y, Hinton G. Deep learning. *Nature*. 2015;521(7553):436.
20. Oikarinen T, Srinivasan K, Meisner O, Hyman JB, Parmar S, Fanucci-Kiss A, et al. Deep convolutional network for animal sound classification and source attribution using dual audio recordings. *The Journal of the Acoustical Society of America*. 2019;145(2):654-62.
21. Dou X, Shirahata S, Sugimoto H. Functional clustering of mouse ultrasonic vocalization data. *PloS One*. 2018;13(5):e0196834.
22. Friedman J, Hastie T, Tibshirani R. *The elements of statistical learning*. New York, NY, USA:: Springer Series in Statistics; 2001.
23. Marafioti A, Perraudin N, Holighaus N, Majdak P. A context encoder for audio inpainting. *arXiv preprint arXiv:12138*. 2018.
24. Goudbeek M, Cutler A, Smits R. Supervised and unsupervised learning of multidimensionally varying non-native speech categories. *Speech Communication*. 2008;50(2):109-25.

25. Guerra L, McGarry LM, Robles V, Bielza C, Larranaga P, Yuste R. Comparison between supervised and unsupervised classifications of neuronal cell types: a case study. *Developmental Neurobiology*. 2011;71(1):71-82.
26. Musolf K, Hoffmann F, Penn DJ. Ultrasonic courtship vocalizations in wild house mice, *Mus musculus musculus*. *Animal Behaviour*. 2010;79(3):757-64.
27. Wang H, Liang S, Burgdorf J, Wess J, Yeomans J. Ultrasonic vocalizations induced by sex and amphetamine in M2, M4, M5 muscarinic and D2 dopamine receptor knockout mice. *PLOS One*. 2008;3(4):e1893.
28. Oppenheim AV. *Discrete-time signal processing*; Pearson Education India; 1999.
29. Rhodes JD, Alseayab S. The generalized chebyshev low - pass prototype filter. *International Journal of Circuit Theory and Applications*. 1980;8(2):113-25.
30. De Boer E, De Jongh H. On cochlear encoding: Potentialities and limitations of the reverse - correlation technique. *The Journal of the Acoustical Society of America*. 1978;63(1):115-35.
31. Balazs P, Holighaus N, Necciari T, Stoeva D. Frame theory for signal processing in psychoacoustics. *Excursions in Harmonic Analysis*. 5: Springer; 2017. p. 225-68.
32. Chen C-P, Bilmes JA, Kirchhoff K, editors. Low-resource noise-robust feature post-processing on Aurora 2.0. *Seventh International Conference on Spoken Language Processing*; 2002.
33. Fukushima K. Neocognitron: A self-organizing neural network for a mechanism of pattern recognition unaffected by shift in position. *Biological Cybernetics*. 1980; 36(4):193–202.
34. LeCun Y, Bottou L, Bengio Y, Haffner P. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*. 1998;86(11):2278-324.
35. Deng L. The MNIST database of handwritten digit images for machine learning research. *IEEE Signal Processing Magazine*. 2012;29(6):141-2.
36. Krizhevsky A, Sutskever I, Hinton GE, editors. *Imagenet classification with deep convolutional neural networks*. *Advances in Neural Information Processing Systems*; 2012.
37. Kahl S, Wilhelm-Stein T, Hussein H, Klinck H, Kowanko D, Ritter M, et al. Large-scale bird sound classification using convolutional neural networks. *Working notes of CLEF*. 2017.
38. Piczak KJ, editor *Environmental sound classification with convolutional neural networks*. *IEEE 25th International Workshop on Machine Learning for Signal Processing (MLSP)*; 2015: IEEE.
39. Zhang Y-J, Huang J-F, Gong N, Ling Z-H, Hu Y. Automatic detection and classification of marmoset vocalizations using deep and recurrent neural networks. *The Journal of the Acoustical Society of America*. 2018;144(1):478-87.
40. McLoughlin I, Zhang H, Xie Z, Song Y, Xiao W. Robust sound event classification using deep neural networks. *IEEE/ACM Transactions on Audio, Speech, Language Processing*. 2015;23(3):540-52.
41. Jaitly N, Nguyen P, Senior A, Vanhoucke V, editors. *Application of pretrained deep neural networks to large vocabulary speech recognition*. *Thirteenth Annual Conference of the International Speech Communication Association*; 2012.
42. Tjandra A, Sakti S, Neubig G, Toda T, Adriani M, Nakamura S, editors. *Combination of two-dimensional cochleogram and spectrogram features for deep learning-based ASR*. *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* 2015: IEEE.
43. Srivastava N, Hinton G, Krizhevsky A, Sutskever I, Salakhutdinov R. Dropout: a simple way to prevent neural networks from overfitting. *The Journal of Machine Learning Research*. 2014;15(1):1929-58.
44. He K, Zhang X, Ren S, Sun J, editors. *Delving deep into rectifiers: Surpassing human-level performance on imagenet classification*. *Proceedings of the IEEE international conference on computer vision*; 2015.
45. Boureau Y-l, Cun YL, editors. *Sparse feature learning for deep belief networks*. *Advances in neural information processing systems*; 2008.
46. Kingma DP, Ba J. Adam: A method for stochastic optimization. *International Conference on Learning Representation (ICLR)*. 2014.
47. Elkan C, editor *The foundations of cost-sensitive learning*. *International joint conference on artificial intelligence*; 2001: Lawrence Erlbaum Associates Ltd.
48. Maaten Lvd, Hinton G. Visualizing data using t-SNE. *Journal of machine learning research*. 2008;9(Nov):2579-605.
49. Arriaga G, Jarvis ED. Mouse vocal communication system: Are ultrasounds learned or innate? *Brain and Language*. 2013;124(1):96-116.
50. Hanley JA, McNeil BJ. The meaning and use of the area under a receiver operating characteristic (ROC) curve. *Radiology*. 1982;143(1):29-36.