

Detection of Glottal Closure Instants from Raw Speech using Convolutional Neural Networks

Mohit Goyal*, Varun Srivastava*, Prathosh AP

Indian Institute of Technology, Delhi

{goyal.mohit999, varunsrivastava.v, prathoshap}@gmail.com

Abstract

Glottal Closure Instants (GCIs) correspond to the temporal locations of significant excitation to the vocal tract occurring during the production of voiced speech. GCI detection from speech signals is a well-studied problem given its importance in speech processing. Most of the existing approaches for GCI detection adopt a two-stage approach (i) Transformation of speech signal into a representative signal where GCIs are localized better, (ii) extraction of GCIs using the representative signal obtained in first stage. The former stage is accomplished using signal processing techniques based on the principles of speech production and the latter with heuristic-algorithms such as dynamic-programming and peak-picking. These methods are thus task-specific and rely on the methods used for representative signal extraction. However in this paper, we formulate the GCI detection problem from a representation learning perspective where appropriate representation is implicitly learned from the raw-speech data samples. Specifically, GCI detection is cast as a supervised multi-task learning problem solved using a deep convolutional neural network jointly optimizing a classification and regression cost. The learning capability is demonstrated with several experiments on standard datasets. The results compare well with the state-of-the-art algorithms while performing better in the case of presence of real-world non-stationary noise.

Index Terms: GCI detection, epoch extraction, dilated convolutional neural networks, multi-task learning.

1. Introduction

1.1. Background and Previous work

Production of voiced speech is accompanied with sustained oscillations of the vocal folds [1] resulting in a quasi-periodic flow of air-pulses which constitutes the excitation signal to the vocal tract [2]. The instant of significant excitation (within each period) is termed as Epoch which coincides with instant of closure of the glottis [3]. The problem of detecting the precise locations of such Glottal Closure Instants (GCIs) from speech signal has been studied for decades given its importance in several speech processing tasks [4, 5, 6, 7, 8, 9, 10, 11, 12]. Most of the successful GCI detectors adopt a two-stage approach - (i) Obtaining an intermediate representation from the speech signal, which explicitly manifests GCIs as discontinuities, impulses, extremas or as other perceptual events and (ii) detecting precise temporal location of glottal closures using custom-made algorithms. The former stage is based on the observation that GCIs are not comprehensible in the raw speech-signal

domain but exhibit better localization in some other domain (Figure 1 in [9]). Many algorithms rely on the source-filter model for speech production [1] and use signal-processing techniques to estimate a correlate of the source-signal, in which GCIs are better manifested. For instance, [5, 6, 10, 13] choose either linear-prediction residual or glottal flow derivative as the representative signal. Other class of algorithms do not explicitly make any model assumption for speech production rather indirectly use the properties of excitation signal (such as its impulsive nature) and estimate appropriate representations (E.g., zero-frequency filtered signal [7], mean-based signal [8], wavelet-decompositions [14], singularity exponents [11]). During the second stage, aforementioned algorithms employ several heuristics to extract (or refine) the GCIs. These include dynamic programming [5, 6], peak-picking [9, 15] and optimization with regularity constraints [11]. All these methods perform reasonably well albeit they depend largely on their choice of representative signals.

1.2. Context and Scope

With the recent advances made in the area of data-driven representation learning [16], it is possible to operate directly in the signal space and let the learning machinery obtain the appropriate representations given any task and data. This approach has found tremendous success in multiple domains with image, text and audio data [17]. Specifically, convolutional neural networks (CNN) have found their utility in a range of speech processing tasks such as phoneme recognition [18, 19], feature/front-end learning for LVCSR [20, 21, 22], voice-activity detection [23], spoofing detection [24], emotion recognition [25, 26] and speaker identification [27]. The underlying theme in all these methods is to directly operate on the raw speech signal and let the CNN learn the abstract representations for the given task. All these papers demonstrate that features learned from data using CNNs lead to superior performance as compared to hand-crafted feature learning.

Now, data driven algorithms have already been employed on the problem of GCI detection with a significant improvement in localization error and detection accuracy. Several standard machine learning methods such as extremely randomized trees, SVM, k-nearest neighbours and multilayer perceptron over hand engineered features obtained from speech signal have been used for GCI detection [28]. CNNs are also shown to be much more proficient in extracting epochs than non-data-driven methods on pathological acoustic speech acquired from vocally impaired patients [29]. Motivated by the above observations, in this paper, we approach the problem of end-to-end GCI detection from a deep-learning perspective. The key difference in the present formulation as compared to the previous works which uses CNNs for various speech processing tasks is that, by-and-large the scope of previous works is classification of an

* Equal Contribution

The code used for implementation including architecture and hyperparameter settings is available at <https://github.com/VarunSrivastavaIITD/DCNN>.

utterance or segment of speech into classes (phonemes, emotions, speakers) while in the proposed work the interest is in detecting a temporal event (GCI) in the signal. This objective is met by formulating the problem in a novel joint classification-regression framework wherein a temporal event (GCI in this case) is simultaneously detected and localized in a frame. Various experiments are performed on multiple datasets comparing with four state-of-the-art algorithms to demonstrate effectiveness of the proposed method through improved detection and localization metrics.

2. Methodology

2.1. Problem formulation and Data generation

In this work, the problem of GCI detection is formulated using a block-processing approach. It is known that GCI is a temporal event that occurs atmost once in every pitch-period of the voiced speech that could range between 2 milli-seconds (ms) to 20 ms [30]. Based on this observation, we define a detection speech window (w_d) as any speech frame of length equal to 2 ms (minimum possible pitch period). Note that, by definition, every detection window can have atmost one GCI within it. Given a w_d , the primary task is to detect whether or not there is GCI within it, which is formulated as a binary classification problem. Having detected a GCI within a w_d , the next step is to localize the GCI which is cast as a regression task, estimating the distance between the onset of w_d and the location of GCI. Since a 2 ms window comprises too less input features (32 samples at 16kHz) for meaningful feature learning via a deep CNN, a symmetric context of 1.5 ms at either sides of the detection window is considered to generate the input frame w_i . Hence in summary, every input data-sample (w_i) will be of length 5 ms, with the classification and regression for GCI detection and localization being carried simultaneously. Finally, all input data-frames of size w_i are generated with a shift of one sample between successive input frames. In principle, heuristics can be designed to reject data-frames (with no associated ground-truth GCI) based on frequency components, amplitude or other signal attributes. However, we refrain from performing any pre-processing/filtering on the input speech to solve the problem in an end to end fashion. Therefore, each possible data-frame is considered for the purpose of epoch extraction.

2.2. Dilated CNNs and Multi-task learning

Recently, Deep Dilated Convolutional Neural Networks have been shown to learn useful representations for several speech processing tasks [31][32]. Since, a dilated convolution is mathematically equivalent to regular convolution with a kernel with zeros inserted in between, stacked dilated convolutions allow an exponential increase in context with a linear increase in parameters. The resulting reduction in parameters provides a regularization effect to the CNN improving its generalization on test data. The proposed neural architecture comprises of convolutional layers that function as feature extractors, so that a *shared* representation is obtained for a speech window. Max pooling is used after every convolutional layer with a kernel size and stride of 2 while doubling the dilation at every convolutional layer to increase the corresponding receptive field at each layer. Scaled Exponential Linear Units (SELU) [33] activation functions are used between convolutional layers in order to normalize the activations. SELU ensures the benefits of explicit normalization (E.g., batch-normalization [34]) while being identical implementation-wise during train and test phase,

unlike techniques such as batch normalization.

Since both the classification and regression tasks are to be performed on the same input frame (w_i) the features are jointly learned which is the core idea behind multitask learning [35]. Thus, an input frame w_i is first fed to the CNN to extract the feature vector that goes into two identical subnetworks performing classification and regression. A sigmoid activation is used at the output of the classification branch and a Hard-tanh function is used for regression (instead of a linear activation to suppress out-of-window predictions).

The task is to output the probability of presence of a GCI y_c ($0 \leq y_c \leq 1$), as well as its location y_r ($0 \leq y_r < w_d$), within each detection window. Ideally, for a window with a GCI, true $y_c = 1$ and true $y_r = d$ (where d is the distance between the onset of the window and location of the GCI within the window) and for a window without a GCI, true $y_c = 0$, and true y_r can be anything and is immaterial (Refer eq.(1)). Since a single network predicts both the probability of occurrence of the GCI and location of the GCI, the loss function consists of two terms, the classification error which is standard cross entropy loss and the regression error which is the mean squared loss between the actual and the predicted location of the GCI within a window. Mathematically,

$$\mathcal{L} = \frac{c_1}{N} \sum_{j=1}^N (t_c^j \log(y_c^j) + (1 - t_c^j) \log(1 - y_c^j)) + \frac{c_2}{\sum_{k=1}^N t_c^k} \sum_{j=1}^N t_c^j * (t_r^j - y_r^j)^2 \quad (1)$$

where t_c^j is 1 if the j^{th} sample window has a GCI, and is 0 otherwise. t_r^j is the actual location of the GCI within the j^{th} window, N is the number of datapoints within a training batch. c_1 and c_2 are the weights of the classification and regression loss respectively, found empirically using grid search. Note that the regression loss is divided by the number of true GCIs in a batch (instead of total number of datapoints) to account for the fact that there are fewer samples for the regression block in every batch compared to the classification block.

2.3. Inference and Weighted Histogram Clustering (WHC)

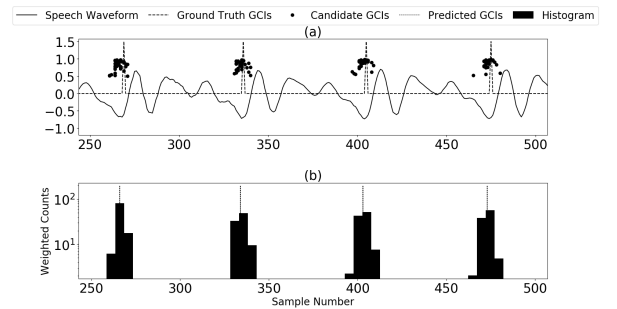


Figure 1: Illustration for WHC (a) a segment of speech with four ground truth GCIs along with the corresponding candidate GCIs (height represents confidence in detection), (b) weighted histograms constructed from the candidate GCIs along with the means of the local-groups (dotted lines) which form the final predictions.

In the present formulation, every input data-frame/window contains atmost one GCI. Since one ground truth GCI can be-

long to multiple data-frames, possibly, there can be duplicate predictions over successive input windows resulting in false detections. Hence, there is a need for multiple predictions to be merged together into a single predicted GCI. Given a speech signal and its associated candidate GCIs, the first step is to construct a weighted histogram of the candidate GCIs with a bin size of B (hyper-parameter). The confidence associated with each candidate for the presence of GCI is used to calculate the weight for the histogram. Since occurrence of glottal closure is a quasi-periodic event, the histogram thus obtained will contain repeating local-groups of contiguous bins with non-zero values (Fig. 1(b)). The final GCIs are hypothesized to be the means (measure of central tendency) of local-group thus (dotted lines in Fig. 1(b)) obtained. These means tend to be at the locations with large number of high-probability candidate GCIs which is desirable. The bin-size B is set at 5 samples through-out this work albeit there is a trade-off between erroneous detections (higher the B lesser are the false insertions) and resolution of predicted GCIs (lower the B , better is the localization of detected GCIs). The entire procedure for clustering is illustrated in Figure 1. The complete algorithm is depicted in Figure 2.

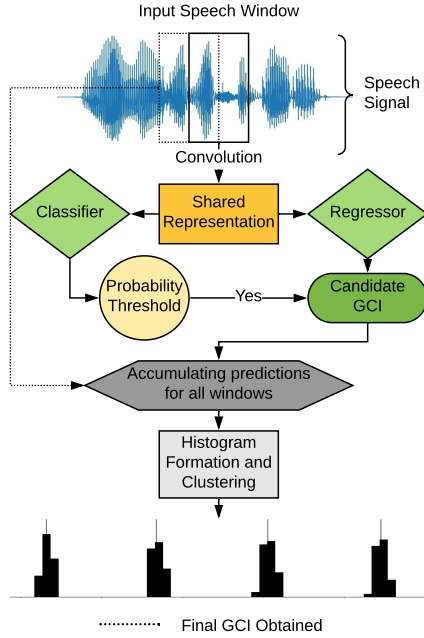


Figure 2: *Proposed Algorithm: Trained CNN outputs a probability and location for each window. The windows satisfying probability threshold are considered as GCI candidates. All candidates are combined and the WHC algorithm gives out the GCIs.*

3. Experiments and Results

3.1. Experimental details

The proposed GCI detection algorithm is evaluated on datasets containing simultaneous recordings of speech and Electroglottography (EGG) signals. The negative peaks obtained from differentiated EGG (dEGG) signals form the ground truth for GCI locations. In this paper, we use speech signals from the CMU Arctic [36] dataset for our experiments. Standard performance

metrics namely, Identification Rate (IDR - % of correct detections, higher the better), Miss Rate (MR - % of missed detections, lower the better), False Alarm Rate (FAR - % of false insertions, lower the better) and Identification Accuracy (IDA - standard deviation of distance between the true and predicted GCIs, lower the better) are employed for evaluation. (A detailed description for metrics may be found in Figure 2 of [5]). There are in total of 3 speakers in the CMU dataset—JMK (Canadian Male), BDL (US Male), SLT (US Female). The CNN model is implemented using PyTorch [37] using the ADAMAX optimizer [38] with standard parameter settings. Model parameters were initialized as per the scheme outlined in [33]. Detection experiments were carried out on clean speech as well as speech corrupted with additive synthetic white noise and real-world babble noise (background multi-speaker chatter obtained from [39]) with SNRs ranging from 0 to 25 dB in steps of 5 dB. For a given SNR, both the training and testing are carried on the corresponding corrupted speech at the same SNR. The results of the proposed algorithm (CNN) are compared with four state-of-the-art algorithms, namely, Zero-frequency resonator [7], Speech Event Detection using the Residual Excitation And a Mean-based Signal (SEDREAMS) [8], Dynamic Plosion Index (DPI) [9] and Micro-canonical Multi-scale Formalism (MMF) [11], by evaluating all four on the same test dataset as that of CNN. Training and testing is performed on CMU corpus with speaker-overlap between the training and test data. Both training and test data contain utterances (non-overlapping across training and test) from all speakers. We also compare the results of extremely randomized trees (ERT) (introduced in [28] for GCI detection) with the proposed algorithm. Since ERT is originally proposed only for clean speech, the comparison is limited to noise free speech. In all experiments a random subset of 10% of data is considered for training and rest for testing. Since 90% of the data is used for testing, no cross-validation is employed.

3.2. Results and Discussion

Table 1 describes the results of six algorithms on the CMU ARCTIC on clean speech. It can be seen that CNN has the highest detection rate (99.3 % and 95.5 %) on both the datasets. The IDA of the CNN method is also the least for both the datasets.

Metrics	IDR(%)	MR(%)	FAR(%)	IDA(ms)
DCNN	99.3	0.3	0.4	0.2
ZFR	96.24	2.7	1.0	0.47
SEDREAMS	92.6	7.4	0.4	0.8
DPI	98.9	0.4	0.3	0.4
MMF	96.0	2.1	1.7	0.5
ERT	93.57	2.16	4.26	0.28

Table 1: *Results on clean speech from CMU Arctic (BDL, JMK, SLT averaged) on six GCI detectors.*

A similar trend is observed even with noisy speech as can be observed in Figure 3. It is seen that CNN and ZFR are very robust to noise with the IDR being above 90% even at 0 dB babble noise. Further, the IDA of the CNN algorithm is consistently lowest for all cases considered.

It is noteworthy that in all these experiments, the model parameters (detection probability threshold: 0.5 and bin-size: 5) were kept the same in spite of which the models generalize well across datasets, recording settings, speakers and noise-levels. However, in practical settings, one can tweak the model

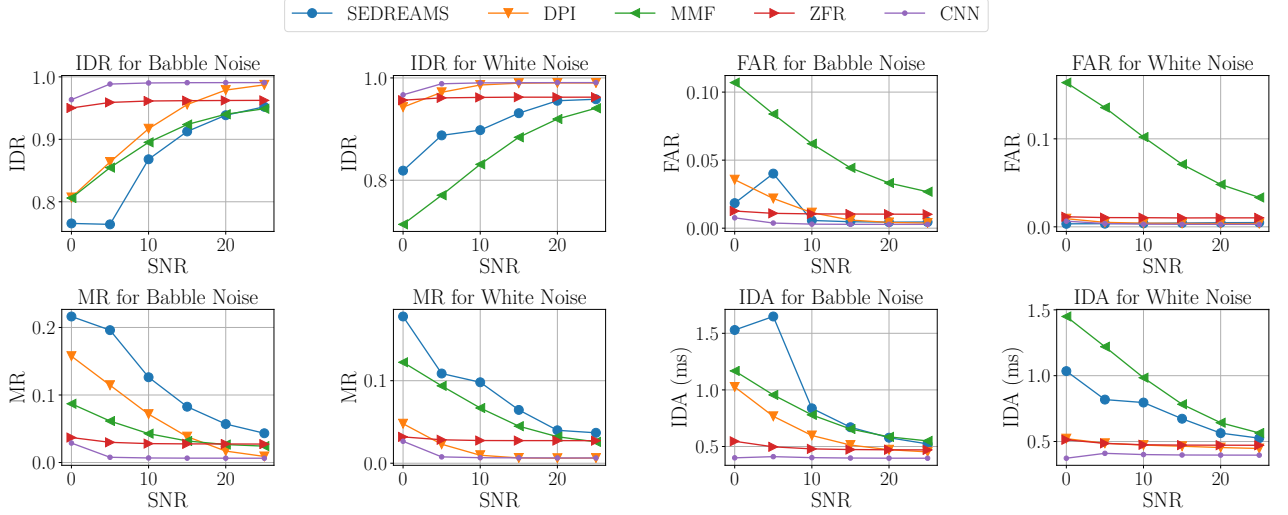


Figure 3: Comparison of five algorithms for different noise levels. It can be seen that the proposed algorithm (CNN) is comparable to the best of the state-of-the-art methods for all cases considered.

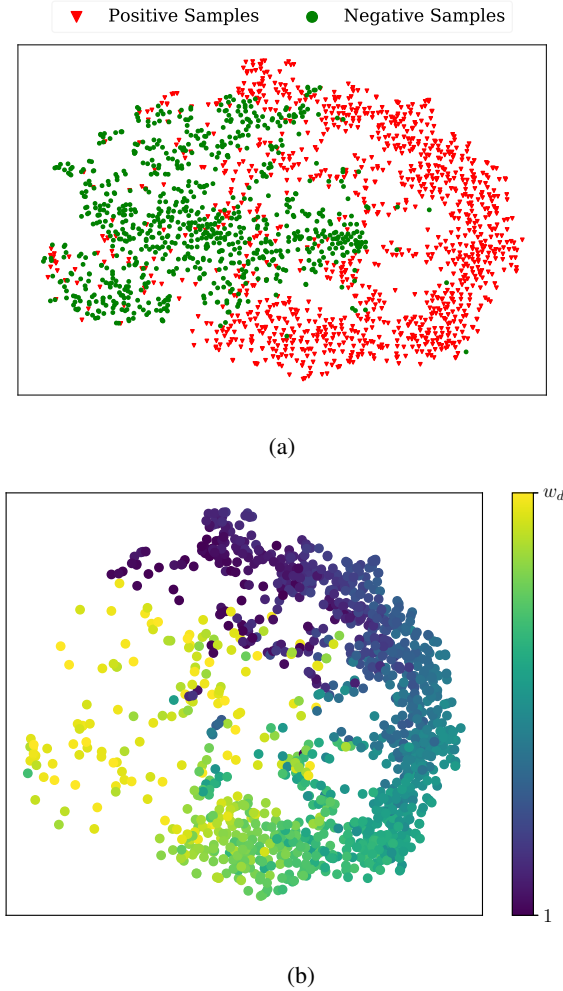


Figure 4: Visualizing shared representation (extracted from CNN) with the help of t-SNE using (a) classification labels. (b) regression values corresponding to color gradient.

hyper-parameters to best suit the data. Given that the proposed method operates on raw speech and fully data-driven, the aforementioned performance is significant. All the results confirm the learning abilities of the proposed architecture. The performance of DCNN might be ascribed to the following factors: (i) problem formulation that enforces presence of multiple detections per GCI leading to robustness, (ii) a multi-task approach to learn a shared feature representation, (iii) use of weighted histogram for clustering.

We also visualize the representations learnt by the shared convolutional layers using t-Distributed Stochastic Neighbor Embedding (t-SNE) [40] technique for dimensionality reduction. It can be seen that the data-points are easily segregated on the basis of class (i.e. whether a given window contains a GCI) as shown in Figure 4a showing that the learnt representation is conducive for high classification accuracy. Simultaneously, we can also observe in Figure 4b that the positive samples corresponding to windows containing GCIs follow a well defined manifold according to different positions of the GCI within a window. Hence, these representations provide a justification for using a multi-task learning paradigm which learns efficient representations suitable for multiple downstream tasks (i.e. classification and regression).

4. Conclusion

In this paper, a data-driven method for GCI detection from raw speech using a multi-task supervised learning approach is proposed. A CNN is employed for learning a shared representation for two tasks simultaneously followed by a distribution estimation method for inference and clustering. Several experiments were conducted to compare the performance with state-of-the-art algorithms on data-sets comprising multiple speakers to demonstrate the efficacy of the proposed approach. Given the representation and generalization abilities of the proposed approach, we believe that a similar methodology could be adopted for detecting multiple landmarks occurring in speech and general time-series data, which could provide directions for future work.

5. References

- [1] G. Fant, *Acoustic theory of speech production: with calculations based on X-ray studies of Russian articulations*. Walter de Gruyter, 1971, vol. 2.
- [2] B. H. Story, "An overview of the physiology, physics and modeling of the sound source for vowels," *Acoustical Science and Technology*, vol. 23, no. 4, pp. 195–206, 2002.
- [3] T. V. Ananthapadmanabha and B. Yegnanarayana, "Epoch extraction from linear prediction residual for identification of closed glottis interval," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. 27, no. 7, pp. 309–319, Aug. 1979.
- [4] T. Ananthapadmanabha and B. Yegnanarayana, "Epoch extraction of voiced speech," *IEEE Transactions on Acoustics, Speech, and Signal Processing*, vol. 23, no. 6, pp. 562–570, December 1975.
- [5] P. A. Naylor, A. Kounoudes, J. Gudnason, and M. Brookes, "Estimation of glottal closure instants in voiced speech using the DYPSA algorithm," *IEEE Trans. Audio, Speech Lang. Process.*, vol. 15, no. 1, pp. 34–43, Jan. 2007.
- [6] M. R. P. Thomas, J. Gudnason, and P. A. Naylor, "Estimation of glottal opening and closing instants in voiced speech using the YAGA algorithm," *IEEE Trans. Audio, Speech, Lang. Process.*, vol. 20, no. 1, pp. 82–91, Jan. 2012.
- [7] K. S. R. Murty and B. Yegnanarayana, "Epoch extraction from speech signals," *IEEE Trans. Audio, Speech, Lang. Process.*, vol. 16, no. 8, pp. 1602–1613, Nov. 2008.
- [8] T. Drugman and T. Dutoit, "Glottal closure and opening instant detection from speech signals," in *Proc. Interspeech Conf.*, 2009.
- [9] A. P. Prathosh, T. V. Ananthapadmanabha, and A. G. Ramakrishnan, "Epoch extraction based on integrated linear prediction residual using plosion index," *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 21, no. 12, pp. 2471–2480, Dec 2013.
- [10] A. P. Prathosh, P. Sujith, A. G. Ramakrishnan, and P. K. Ghosh, "Cumulative impulse strength for epoch extraction," *IEEE Signal Processing Letters*, vol. 23, no. 4, pp. 424–428, 2016.
- [11] V. Khanagha, K. Daoudi, and H. M. Yahia, "Detection of glottal closure instants based on the microcanonical multiscale formalism," *IEEE/ACM Transactions on Audio, Speech and Language Processing (TASLP)*, vol. 22, no. 12, pp. 1941–1950, 2014.
- [12] T. Drugman, M. Thomas, J. Gudnason, P. Naylor, and T. Dutoit, "Detection of glottal closure instants from speech signals: A quantitative review," *IEEE Trans. Audio, Speech, Lang. Process.*, vol. 20, no. 3, pp. 994–1006, Mar. 2012.
- [13] A. I. Koutrouvelis, G. P. Kafentzis, N. D. Gaubitch, and R. Heusdens, "A fast method for high-resolution voiced/unvoiced detection and glottal closure/opening instant estimation of speech," *IEEE/ACM Transactions on Audio, Speech and Language Processing (TASLP)*, vol. 24, no. 2, pp. 316–328, 2016.
- [14] C. D' Alessandro and N. Sturm, "Glottal closure instant and voice source analysis using time-scale lines of maximum amplitude," *Sadhana*, vol. 36, no. 5, pp. 601–622, 2011.
- [15] T. Drugman, M. Thomas, J. Gudnason, P. Naylor, and T. Dutoit, "Detection of glottal closure instants from speech signals: A quantitative review," *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 20, no. 3, pp. 994–1006, March 2012.
- [16] Y. Bengio, A. Courville, and P. Vincent, "Representation learning: A review and new perspectives," *IEEE transactions on pattern analysis and machine intelligence*, vol. 35, no. 8, pp. 1798–1828, 2013.
- [17] I. Goodfellow, Y. Bengio, A. Courville, and Y. Bengio, *Deep learning*. MIT press Cambridge, 2016, vol. 1.
- [18] D. Palaz, R. Collobert, and M. Magimai-Doss, "Estimating phoneme class conditional probabilities from raw speech signal using convolutional neural networks," in *Proc. Interspeech Conf.*, 2013.
- [19] N. Zeghidour, U. Nicolas, K. Iasonas, S. Thomas, S. Gabriel, and D. Emmanuel, "Learning filterbanks from raw speech for phone recognition," in *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2018.
- [20] T. N. Sainath, R. J. Weiss, A. Senior, K. W. Wilson, and O. Vinyals, "Learning the speech front-end with raw waveform CLDNNs," in *Proc. Interspeech Conf.*, 2015.
- [21] Z. Tuske, P. Golik, R. Schlueter, and H. Ney, "Acoustic modeling with deep neural networks using raw time signal for LVCSR," in *Proc. Interspeech Conf.*, 2014.
- [22] O. Abdel-Hamid, A.-r. Mohamed, H. Jiang, L. Deng, G. Penn, and D. Yu, "Convolutional neural networks for speech recognition," *IEEE/ACM Transactions on audio, speech, and language processing*, vol. 22, no. 10, pp. 1533–1545, 2014.
- [23] R. Zazo, T. N. Sainath, G. Simko, and C. Parada, "Feature learning with raw-waveform CLDNNs for voice activity detection," in *Proc. Interspeech Conf.*, 2016.
- [24] H. Dinkel, N. Chen, Y. Qian, and K. Yu, "End-to-end spoofing detection with raw waveform cldnns," in *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2017, pp. 4860–4864.
- [25] G. Trigeorgis, F. Ringeval, R. Brueckner, E. Marchi, M. A. Nicolaou, B. Schuller, and S. Zafeiriou, "Adieu features? end-to-end speech emotion recognition using a deep convolutional recurrent network," in *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2016, pp. 5200–5204.
- [26] S. Zhang, S. Zhang, T. Huang, and W. Gao, "Speech emotion recognition using deep convolutional neural network and discriminant temporal pyramid matching," *IEEE Transactions on Multimedia*, vol. 20, no. 6, pp. 1576–1590, 2017.
- [27] H. Muckenhirn, M. Magimai-Doss, and S. Marcel, "Towards directly modeling raw speech signal for speaker verification using cnns," in *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2018.
- [28] J. MatouÅek and D. Tihelka, "Classification-based detection of glottal closure instants from speech signals," 08 2017, pp. 3053–3057.
- [29] G. R. M., T. Mandal, and K. S. Rao, "Glottal closure instants detection from pathological acoustic speech signal using deep learning," *Machine Learning for Health (ML4H) Workshop at NeurIPS*, 2018. [Online]. Available: <http://arxiv.org/abs/1811.09956>
- [30] I. Titze, *Principles of Voice Production*, 1994, p. 188.
- [31] F. Yu and V. Koltun, "Multi-scale context aggregation by dilated convolutions," *arXiv preprint arXiv:1511.07122v3*, 2015.
- [32] A. Van Den Oord, S. Dieleman, H. Zen, K. Simonyan, O. Vinyals, A. Graves, N. Kalchbrenner, A. Senior, and K. Kavukcuoglu, "Wavenet: A generative model for raw audio," *arXiv preprint arXiv:1609.03499v2*, 2016.
- [33] G. Klambauer, T. Unterthiner, A. Mayr, and S. Hochreiter, "Self-normalizing neural networks," in *Advances in Neural Information Processing Systems*, 2017, pp. 972–981.
- [34] S. Ioffe and C. Szegedy, "Batch normalization: Accelerating deep network training by reducing internal covariate shift," *arXiv preprint arXiv:1502.03167*, 2015.
- [35] R. Caruana, "Multitask learning," in *Learning to learn*. Springer, 1998, pp. 95–133.
- [36] J. Kominek and A. W. Black, "The CMU arctic speech databases," in *Fifth ISCA Workshop on Speech Synthesis*, 2004.
- [37] P. C. Team, "Pytorch: Tensors and dynamic neural networks in python with strong GPU acceleration," 2017.
- [38] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," *arXiv preprint arXiv:1412.6980v9*, 2014.
- [39] NoiseX-92. [Online]. Available: www.speech.cs.cmu.edu/comp.speech/SectionI/Data/noisex.html
- [40] L. v. d. Maaten and G. Hinton, "Visualizing data using t-sne," *Journal of machine learning research*, vol. 9, no. Nov, pp. 2579–2605, 2008.