

SPEECH RECOGNITION WITH NO SPEECH OR WITH NOISY SPEECH

Gautam Krishna*

Co Tran*[‡]

Jianguo Yu^{†‡}

Ahmed H Tewfik*

*Brain Machine Interface Lab, The University of Texas at Austin

[†] University of Aizu

ABSTRACT

The performance of automatic speech recognition systems(ASR) degrades in the presence of noisy speech. This paper demonstrates that using electroencephalography (EEG) can help automatic speech recognition systems overcome performance loss in the presence of noise. The paper also shows that distillation training of automatic speech recognition systems using EEG features will increase their performance. Finally, we demonstrate the ability to recognize words from EEG with no speech signal on a limited English vocabulary with high accuracy.

Index Terms— Electroencephalography(EEG), Speech Recognition, Distillation, Deep Learning

1. INTRODUCTION

Traditional state of art Automatic Speech Recognition (ASR) systems mainly uses acoustic features for doing speech recognition. In [1] authors show how combining acoustic and articulatory features can help in designing robust speech recognition systems. Recently, researchers have used Functional near-infrared spectroscopy (fNIRS) signals for doing speech recognition with 74.7 % accuracy [2]. In [3, 4] authors provide interesting results on how Electroencephalography (EEG) signals, which is an invasive approach can help in speech recognition.

Electroencephalography (EEG) on other hand is a non invasive approach. It is a measure of electrical activity of the human brain. In [5] authors demonstrate decoding vowel articulation using EEG cortical currents. In our work we used only surface EEG potentials, which are directly obtained from the EEG sensors.

In [6] authors used EEG signals to perform envisioned speech recognition using random forest algorithm and they reported an average accuracy of 85.2 %. In our work we used a deep learning model and achieved a highest test accuracy of 99.38 %. In [7] authors propose neural network based model which predicts phonemes from EEG but in our work the model directly predicts words with higher accuracy and we also study the effect of noisy speech. References [8, 9, 10, 11] describes various techniques to perform speech recognition with noisy

speech but as far as we know our work is first demonstration of EEG's ability to overcome performance loss in presence of background noise and our approach demonstrates a significant high recognition accuracy of 99.38 % for recognition of limited words in presence of background noise.

In [12] Hinton proposed the concept of distilling the knowledge in neural networks, where a simple model learns a complex task by imitating the solution of a more complicated and flexible model. In [13] authors demonstrate the integration of acoustic and articulatory features using Generalized Distillation.

Motivated by the results presented in [2, 7, 14, 5, 15, 6, 16, 17], primary goal of our research was to create a state of art ASR system and train it with EEG features, acoustic features, concatenation of acoustic - EEG features and investigate its performance in absence and presence of background noise. We tested our model for recognition of the five English vowels and four English words 'yes', 'no', 'left', 'right'.

Inspired from [12] and [13] we further tried implementing Generalized Distillation in the speech recognition task in order to integrate EEG information into speech recognition system and observed that distillation training with EEG improves the recognition accuracy of our ASR model.

Major contributions of our paper are as follows: we identified a set of EEG features which are better representation of speech, we proposed a deep learning model that is capable of learning EEG features and perform speech recognition with no speech as input, we demonstrate the ability of EEG features to make up for ASR performance lost due to background noise and we show that performance of ASR system can be improved by distillation training with EEG features.

2. AUTOMATIC SPEECH RECOGNITION SYSTEM MODEL

For this work we created an ASR model using gated recurrent unit (GRU) networks [18]. The model was created using Google TensorFlow deep learning library. GRU has an architecture similar to long short term memory (LSTM) but has less parameter's compared to LSTMs. Hence GRU's are ideal for recurrent neural network applications where less amount of training data is available. A GRU cell architecture is shown

[‡] Equal author contribution

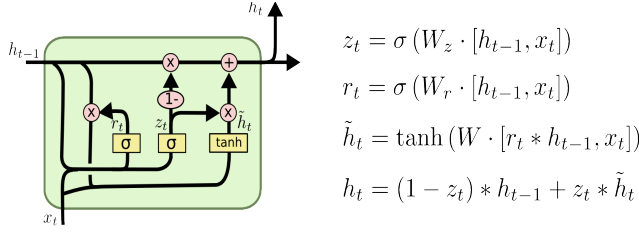


Fig. 1. Gated Recurrent Unit Cell

in Figure 1, x_t is the input vector, h_t is the output vector, z_t is update gate vector and r_t is reset gate vector. Sigmoid and hyperbolic tangent activation functions are used in the GRU cell.

Given the sequence of input vectors $X = [x_0, x_1, \dots, x_n]$ and $x_i \in R^d$ is the input vector at time i . Our model contains a GRU layer, an average pooling layer and a dense (fully connected) layer followed by output layer, which takes X and returns the probabilities being the all words (vowels) $o \in R^p$, as shown in Figure 2. The pooling layer computes the average value $\frac{1}{n} \sum_{i=0}^n h_i$ of all outputs of GRU layer.

The number of hidden units in GRU layer is 128 and in dense layer is 64 respectively, the output dimension p is 4 or 5 corresponding to the number of classes. The batch size is 1 and the dropout rate for the dense layer is 0.2. We used Adam Optimizer with 0.001 learning rate. We used cross entropy as the loss function.

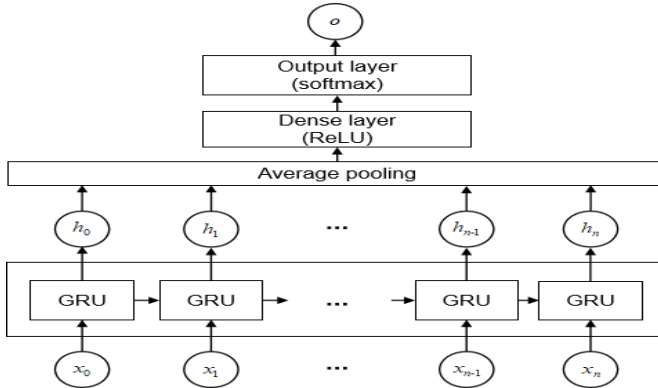


Fig. 2. Our ASR Model Architecture

As shown in Table 1, 64 percent of the data was used for training set, 16 percent for validation set and remaining 20 percent for testing set.

We trained and tested GRU based deep learning ASR model using three different feature sets. 1) only acoustic features, 2) concatenation of acoustic and EEG features and 3) only EEG features. Number of training epochs was 10000 for all cases except for vowels in presence of noise data set. For that data set the number of epochs was set to 30000 as we didn't observe convergence at 10000 epochs. In Figure 2 observation vectors x_0, x_1 up to x_n are treated as input vectors.

Words/Vowel	Class	Training set	Validation set	Test set	Total
	Ratio	64	16	20	100
Words	yes	195	49	61	305
Words	no	259	66	81	406
Words	right	219	56	68	343
Words	left	214	54	67	335
Vowel	a	170	44	53	267
Vowel	e	170	44	53	267
Vowel	i	170	44	53	267
Vowel	o	170	44	53	267
Vowel	u	170	44	53	267

Table 1. Training, Validation and Test sets

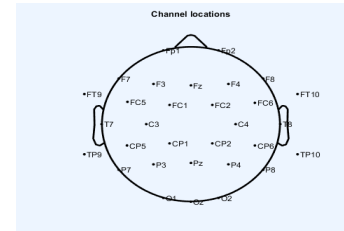


Fig. 3. EEG channel locations for the cap used in our experiments

The index denotes the time step value. There was no fixed time step value for the GRU. We used dynamic recurrent neural network (RNN) cell of tensorflow. Observation vector X can be Mel-frequency cepstral coefficients (MFCC) / acoustic features, EEG or concatenation of EEG and MFCC features depending on how the ASR model was trained.

3. GENERALIZED DISTILLATION

Distillation training involves following three main steps.

1. Train an ASR model with EEG + MFCC. This model is called the teacher model.
2. Generate soft targets from this model. After training the teacher model, we used `estimator.predict` from tensorflow (temperature parameter) to compute the unscaled prediction probabilities for each training example. This unscaled prediction probabilities are called soft targets.
3. Train an ASR model with MFCC + soft targets. This model is called the student model.

The hyper parameters to tune for the training the student model are temperature and lambda. Temperature is a hyper parameter of neural network used to control the randomness of predictions by scaling the logits (raw predictions) before applying softmax activation function [12]. Training loss for

the student model is defined as:

$$\sum (hard_loss * (1 - \lambda) + soft_loss * \lambda) \quad (1)$$

Where *soft_loss* is the cross entropy loss between soft targets, soft logits and *hard_loss* is the cross entropy loss between hard targets, logits. The parameter λ behaves like a regularization parameter in the loss function of the student model. It is called the imitation parameter.

We tuned the hyper parameters temperature and λ for the following values:- Temperature = [1.0, 2.0, 5.0, 10.0] and λ = [0, 0.2, 0.8, 1.0] [12]. We used the ASR model shown in Figure 2 for both our teacher and student model. Instead of the average pooling layer, we used the last time step output of the GRU layer as input to the dense layer for distillation training.

4. DESIGN OF EXPERIMENTS FOR BUILDING THE DATABASE

Four subjects took part in the experiment. All subjects were male undergraduate students in their early twenties. Three were native English speakers and one had an accent. The subjects were asked to speak English vowels [a, e, i, o, u] separately for 5 minutes each with a time interval of 2 seconds for each vowel. Their simultaneous speech and EEG signals were recorded. The same subjects were then asked to speak the English words [yes, no, left, right]. Their simultaneous speech and EEG signals were recorded. The words were spoken for 5 minutes each with a time interval of 5 seconds for each word.

We then repeated the same set of EEG-Speech recording experiments for recognition of English words and vowels in presence of background noise. For generating a background noise of 60 db we used background music played from our lab computer as the source of the noise. We used Brain Vision EEG recording hardware. Our EEG cap had 32 wet EEG electrodes including one electrode as **ground** as shown in Figure 3. We used EEGLab [19] to obtain the EEG sensor location mapping. It is based on standard 10-20 EEG sensor placement method for 32 electrodes [20].

In total for this work we used 75 minutes of speech EEG data for vowels with noise, 75 minutes of speech EEG data for vowels without noise, 40 minutes of speech EEG data for words with noise and 40 minutes of speech EEG data for words without noise. The data was recorded from the subjects on different days.

5. EEG AND SPEECH FEATURE EXTRACTION DETAILS

EEG signals were sampled at 1000Hz and a fourth order IIR band pass filter with cut off frequencies 0.1Hz and 70Hz was

applied. A notch filter with cut off frequency 60 Hz was used to remove the power line noise. EEGLab's [19] Independent component analysis (ICA) toolbox was used to remove other biological signal artifacts like electrocardiography (ECG), electromyography (EMG), electrooculography (EOG) etc from the EEG signals. We extracted five statistical features for EEG, namely root mean square, zero crossing rate, moving window average, kurtosis and power spectral entropy [21]. So in total we extracted 31(channels) X 5 or 155 features for EEG signals. The EEG features were extracted at a sampling frequency of 100Hz for each EEG channel.

The recorded speech signal was sampled at 16KHz frequency. We extracted MFCC as features for speech signal. We first extracted MFCC 13 features and then computed first and second order differentials (delta and delta-delta) thus having total MFCC 39 features. The MFCC features were also sampled at 100Hz same as the sampling frequency of EEG features to avoid seq2seq problem.

6. EEG FEATURE DIMENSION REDUCTION ALGORITHM DETAILS

After extracting EEG and acoustic features as explained in the previous section, we used non linear methods to do feature dimension reduction in order to obtain set of EEG features which are better representation of acoustic features. We reduced the 155 EEG features to a dimension of 39 by applying Kernel Principle Component Analysis (KPCA). We plotted cumulative explained variance versus number of components to identify the right feature dimension. We used KPCA with polynomial kernel of degree 3. We further computed delta, delta and delta of those 39 EEG features, thus the final feature dimension of EEG was 117 (39 times 3). This approach gave best performance for feature dimension reduction for EEG data recorded for words in presence, absence of background noise and for vowels in presence of background noise. For EEG data recorded for vowels in absence of background noise we used autoencoder for doing feature dimension reduction. Here the EEG feature dimension was first reduced to 6 by autoencoder and delta, delta and delta features were computed thus making the final EEG feature dimension equal to 18 for that data set.

7. RESULTS

The evaluation metric was recognition accuracy. We defined ASR recognition accuracy as the ratio of number of correct predictions given by our model to the total number of predictions given by the model in training, validation and test data sets respectively.

Table 2 and 3 shows validation accuracy, test accuracy values obtained after convergence for different data sets when trained using acoustic only, concatenation of acoustic and

Words/Vowels	Background noise	MFCC acc	MFCC-EEG acc	EEG acc
Vowels	No	88.75	97.50	91.25
Vowels	Yes	73.33	92.00	92.00
Words	No	95.83	97.91	96.52
Words	Yes	94.53	98.39	98.39

Table 2. Validation accuracy for **ASR EEG fusion** for different datasets

EEG, EEG only features respectively. The test, train and validation accuracy values were comparable indicating that our model didn't over fit.

When we trained the model using 31 EEG channels + MFCC, we obtained a test accuracy of 96.36% on vowels in absence of noise dataset as shown in Table 3. In order to make the system more applicable to real world, we also trained the model with a smaller feature set containing only 4 EEG channels (T7, T8, Fc5, P7) + MFCC and obtained a remarkably close 93% accuracy on the same dataset.

We obtained best results during test time after distillation training for the hyper parameters temperature equal to 2 and lambda equal to 0.2, as shown in Table 4.

Table 5 shows test accuracy values obtained after distillation training for different data sets. Student model underwent distillation training but during its test time EEG features are not provided.

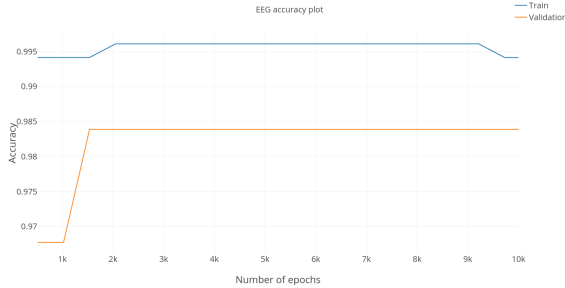


Fig. 4. ASR using EEG only accuracy plot for recognition of words in presence of background noise

8. CONCLUSION

To our knowledge, this is the first time that a deep learning model based speech recognition is demonstrated with high accuracy using only EEG features for character or word level prediction. Our work also demonstrates the ability of EEG to make up for ASR performance lost due to background noise. We also show that distillation training can improve the accuracy of an ASR system fused with EEG features. We are currently working towards speech recognition for a larger speech

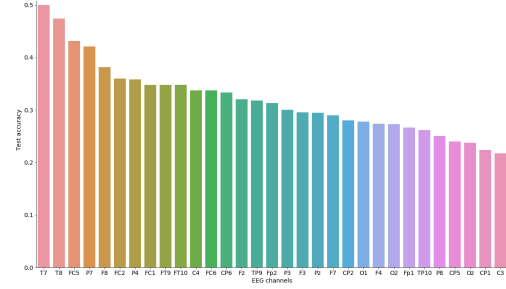


Fig. 5. ASR test accuracy (EEG only) contribution per each EEG sensor. Channels **T7, T8, Fc5** and **P7** showed highest contribution to test accuracy. Dataset used here was vowels in absence of noise.

Words/Vowels	Background noise	MFCC acc	MFCC-EEG acc	EEG acc
Vowels	No	89.09	96.36	90.91
Vowels	Yes	74.74	94.74	92.63
Words	No	95.63	97.91	96.87
Words	Yes	93.00	97.50	99.38

Table 3. Test accuracy for **ASR EEG fusion** for different datasets

Temp	Lambda	Student Training acc	Student Validation acc	Student Test-ing acc
1	0.0	99.39	97.22	97.22
2	0.2	99.54	97.22	98.61
5	0.8	99.23	95.83	98.61
10	1	94.94	95.83	94.44

Table 4. Hyper parameter tuning table for **distillation training**. The data set used here was words with no noise

Words/Vowels	Background noise	Student acc	MFCC acc
Vowels	No	92.73	89.09
Vowels	Yes	76.84	74.74
Words	No	98.61	95.83
Words	Yes	97.62	93.00

Table 5. Test accuracy after **distillation training** for different datasets

EEG corpus. We believe speech recognition using EEG will help people with speaking difficulties to have better technology accessibility.

9. REFERENCES

- [1] Katrin Kirchhoff, Gernot A Fink, and Gerhard Sagerer, "Combining acoustic and articulatory feature information for robust speech recognition," *Speech Communication*, vol. 37, no. 3-4, pp. 303–319, 2002.
- [2] Yichuan Liu and Hasan Ayaz, "Speech recognition via fnirs based brain signals," *Frontiers in Neuroscience*, vol. 12, pp. 695, 2018.
- [3] NF Ramsey, E Salari, EJ Aarnoutse, MJ Vansteensel, MB Bleichner, and ZV Freudenburg, "Decoding spoken phonemes from sensorimotor cortex with high-density ecog grids," *Neuroimage*, 2017.
- [4] Stephanie Martin, Peter Brunner, Iñaki Iturrate, José del R Millán, Gerwin Schalk, Robert T Knight, and Brian N Pasley, "Word pair classification during imagined speech using direct brain recordings," *Scientific reports*, vol. 6, pp. 25803, 2016.
- [5] Natsue Yoshimura, Atsushi Nishimoto, Abdelkader Nasreddine Belkacem, Duk Shin, Hiroyuki Kambara, Takashi Hanakawa, and Yasuharu Koike, "Decoding of covert vowel articulation using electroencephalography cortical currents," *Frontiers in neuroscience*, vol. 10, pp. 175, 2016.
- [6] Pradeep Kumar, Rajkumar Saini, Partha Pratim Roy, Pawan Kumar Sahu, and Debi Prosad Dogra, "Envisioned speech recognition using eeg sensors," *Personal and Ubiquitous Computing*, vol. 22, no. 1, pp. 185–199, 2018.
- [7] Pengfei Sun and Jun Qin, "Neural networks based eeg-speech models," *arXiv preprint arXiv:1612.05369*, 2016.
- [8] Anuroop Sriram, Heewoo Jun, Yashesh Gaur, and Sanjeev Satheesh, "Robust speech recognition using generative adversarial networks," in *2018 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2018, pp. 5639–5643.
- [9] Ian McLoughlin, Haomin Zhang, Zhipeng Xie, Yan Song, and Wei Xiao, "Robust sound event classification using deep neural networks," *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 23, no. 3, pp. 540–552, 2015.
- [10] Jort F Gemmeke, Tuomas Virtanen, and Antti Hurmalainen, "Exemplar-based sparse representations for noise robust automatic speech recognition," *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 19, no. 7, pp. 2067–2080, 2011.
- [11] Tian Tan, Yanmin Qian, Hu Hu, Ying Zhou, Wen Ding, and Kai Yu, "Adaptive very deep convolutional residual network for noise robust speech recognition," *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 26, no. 8, pp. 1393–1405, 2018.
- [12] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean, "Distilling the knowledge in a neural network," 2014.
- [13] Jianguo Yu, Konstantin Markov, and Tomoko Matsui, "Articulatory and spectrum features integration using generalized distillation framework," in *Machine Learning for Signal Processing (MLSP), 2016 IEEE 26th International Workshop on*. IEEE, 2016, pp. 1–6.
- [14] Tanja Schultz, Michael Wand, Thomas Hueber, Dean J Krusienski, Christian Herff, and Jonathan S Brumberg, "Biosignal-based spoken communication: A survey," *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 25, no. 12, pp. 2257–2271, 2017.
- [15] Mashael M AlSaleh, Mahnaz Arvaneh, Heidi Christensen, and Roger K Moore, "Brain-computer interface technology for speech recognition: A review," in *Signal and Information Processing Association Annual Summit and Conference (APSIPA), 2016 Asia-Pacific*. IEEE, 2016, pp. 1–5.
- [16] Marianna Rosinová, Martin Lojka, Ján Staš, and Jozef Juhár, "Voice command recognition using eeg signals," in *ELMAR, 2017 International Symposium*. IEEE, 2017, pp. 153–156.
- [17] Jongin Kim, Suh-Kyung Lee, and Boreom Lee, "Eeg classification in a single-trial basis for vowel speech perception using multivariate empirical mode decomposition," *Journal of neural engineering*, vol. 11, no. 3, pp. 036010, 2014.
- [18] Junyoung Chung, Caglar Gulcehre, KyungHyun Cho, and Yoshua Bengio, "Empirical evaluation of gated recurrent neural networks on sequence modeling," *arXiv preprint arXiv:1412.3555*, 2014.
- [19] Arnaud Delorme and Scott Makeig, "Eeglab: an open source toolbox for analysis of single-trial eeg dynamics including independent component analysis," *Journal of neuroscience methods*, vol. 134, no. 1, pp. 9–21, 2004.
- [20] Frank Sharbrough, "American electroencephalographic society guidelines for standard electrode position nomenclature," *J clin Neurophysiol*, vol. 8, pp. 200–202, 1991.
- [21] Aihua Zhang, Bin Yang, and Ling Huang, "Feature extraction of eeg signals using power spectral entropy," in *2008 International Conference on BioMedical Engineering and Informatics*. IEEE, 2008, pp. 435–439.